



Prediction of Business Bankruptcy with the Help of Extreme Gradient Increase

Catalin-Emanuel CIOBOTA*, Manuela-Violeta TUREATCA**

ARTICLE INFO

Article history:

Accepted December 2022

Available online December 2022

JEL Classification

M41, M42

Keywords:

bankruptcy, enterprise, synthetic minority, insolvency

ABSTRACT

Financial institutions use business failure forecasting models to manage their investments. The accuracy of the forecast is an important factor in determining how much capital is needed to cover credit losses. The majority of studies use traditional statistical methods to model business failures (e.g., discriminant analysis and logistic regression). Models constructed are usually quite inaccurate, however. It is due to the imbalance between the classes of training samples (bankruptcies account for a small percentage of all firms) that are used to construct the models. Currently, various machine learning methods such as the random forest method and the gradient boosting method are widespread. In this study, the main focus is on using extreme gradient growth to predict bankruptcy. Extreme gradient boosting, using LASSO or Ridge regularization, penalizes complex models to avoid overfitting. Also, during training, extreme gradient boosting fills in the missing values in the data set depending on the amount of loss. In this article, in order to increase the efficiency of extreme degree growth, it is proposed to use SMOTE technology to improve class balance. The quality values of the solutions obtained by the improved extremal degree increase are compared with the solutions obtained by other methods.

© 2022 EAI. All rights reserved.

1. Introduction

Investment risks are a major concern for financial institutions, which forces them to verify and control the financial solvency of the enterprise. Bankruptcy models are used by financial institutions to assess the credit losses they may suffer if their counterparties default on their debts. To determine how much capital is needed to cover credit losses, financial institutions need an accurate forecast. Currently, there is a fairly large number of studies aimed at improving the accuracy of corporate bankruptcy models. Practically, the significance of bankruptcy forecasting is related to the significant interest of creditors to accurately predict the future existence of companies.

Most studies treat bankruptcy forecasting as a binary classification problem. The target variable of the models usually takes two values: 0 (solvent company) and 1 (bankrupt company). Bankruptcy estimation, discriminant analysis [1–5] and regression analysis are traditionally used for this. Analysis [5–7]. Altman used linear discriminant analysis (Linear Discriminant Analysis, LDA) to distinguish bankrupt from non-bankrupt enterprises [1].

LDA uses a linear combination of scores to classify a business. This assessment is then applied to separate bankrupt and non-bankrupt enterprises. In [2], LDA is used to predict the bankruptcy risk of enterprises. The applicability of LDA analysis and quadratic discrimination is studied in [4, 5]. In [5], enterprise bankruptcy models are built using logistic regression, but unlike the Altman model [1], the Olson model [5] determines the probability of bankruptcy. Despite the relative ease of building models using discriminant analysis and logistic regression, the use of these models shows poor generalization and low prediction accuracy [8]. Therefore, in the future, many researchers use machine learning methods such as neural networks [9–12], support vector machines [13, 14], amplification methods [15, 16].

One of the characteristic features of the bankruptcy forecasting problem is that the share of bankrupt firms is a few percent of the total number of firms.

Therefore, researchers are faced with unbalanced data sets in which the number of failing firms clearly exceeds the number of non-failing ones. In a series of papers, attempts are made to find methods to improve unbalanced data sets. Most of them use mechanisms to balance the sampling distributions. In [17], several methods are used to improve two highly unbalanced datasets.

*Valahia University of Targoviste, Romania, **Dunarea de Jos University of Galati, Romania. Email addresses: ciobota_catalin@yahoo.com (C. E. Ciobota), mvtdeclaratii@gmail.com (M. V. Tureatca – Corresponding author).

These methods allow bankruptcy forecasting methods to achieve better results than forecasts from the original unbalanced data sets. Paper [18] uses as training a sample containing 620 bankrupt samples and 7398 non-bankrupt samples of firms. We will be using subsampling to improve the accuracy of bankruptcy forecasting. This type of cluster analysis combines nearest neighbors and support vector machines. Research results show that the proposed hybrid method improves the accuracy of classifiers such as logistic regression, discriminant analysis, decision tree, and support vector machines. The paper [19] evaluates four methods for improving the balance of data sets: random oversampling, random undersampling, minority augmentation technique (SMOTE) and EasyEnsemble. Research shows that SMOTE is a more accurate data balance improvement method. The purpose of this study is to investigate the effectiveness of the extreme stimulation method in combination with the training balance improvement method. It offered a comparison of the results obtained with the solutions obtained by other methods.

The article is organized as follows: Section 1 introduces the topic of extreme gradient boosting and SMOTE technology for better balancing data. Section 2 provides a dataset with financial performance data from firms and also labels indicating whether or not that company is bankrupt, or will soon be bankrupt. Lastly, it provides the results of a study on the effectiveness of bankruptcy forecasting models. In conclusion, a brief summary of the results obtained in the paper and directions for further research are provided.

2. Theoretical background

2.1. Extreme grade growth

Extreme Gradient Boosting (XGB) is a development of the gradient boosting method [20, 21]. In supervised learning, the dataset consists of objects with features and labels.

It is necessary to build a model that can predict as accurately as possible the labels for each new object. The stimulus model additively uses features for label prediction

$$\tilde{y}_i = F(x_i) = \sum_{j=1}^N f_j(x_i). \quad (1)$$

To train the feature set, we minimize the loss function

$$L(F) = \sum_{i=1}^n l(y_i, \tilde{y}_i) + \sum_{j=1}^K R(f_j), \quad (2)$$

where $R(f_i)$ is the regularization coefficient that limits the complexity of the models. The loss function $L(f)$ contains K functions as parameters, so it is very difficult to optimize directly. Instead, we optimize the model additively. Let y_i be the prediction of sample i at the third iteration. We will add $f(t)$ to minimize

$$L^{(t)} = \sum_{i=1}^n l(y_i, \tilde{y}_i^{(t-1)} + f_t(x_i)) + R(f_t). \quad (3)$$

This means that $f(t)$ is added in such a way that we improve our model as much as possible for each iteration. We use a second-order approximation that uses the gradient of this intermediate loss function. For this reason, we call his gradient boosting algorithm.

The main differences between extreme degree boosting and gradient boosting are related to regularization and working with missing data. Extreme gradient growth penalizes complex models by applying Ridge or LASSO regularization to reduce the risk of overfitting. Also, during training extreme gradient boosting will use a gap-filling algorithm with values based on the magnitude of the losses.

The accuracy of the solutions obtained using XGB depends on the hyperparameters of the method. The most important of these are maximum tree depth, learning rate, number of trees. Increasing the maximum depth of the tree makes the model more complex and increases the probability of overfitting. A higher learning rate results in each tree making more serious corrections to the solution. As a result, model complications and overfitting are more likely to occur. In addition to increasing the complexity of the model, adding more trees also increases its accuracy.

2.2. Extreme gradient increase

In order to improve the quality of bankruptcy forecasting, it is proposed to use the method of increasing the extreme degree, a method to improve the balance of the training set, when building the models. Same as the other technique, we propose to use Synthetic Minority Oversampling Technique (SMOTE), which is based of the idea of generating a certain number of artificial samples similar to those in the

bankruptcy class without duplicating them. To create a new model find the difference $d = x_b - x_a$, where x_b , x_a are characteristic vectors of "neighboring" samples a and b class of bankrupt enterprises. They are found using the nearest neighbor algorithm. In our case, it is necessary and sufficient to generate a new sample to obtain a set of k nearest neighbors from which the objects will be selected.

Next, the difference vector of the new sample is obtained by multiplying each vector element by a random number in the interval (0, 1). The characteristic vector of the new sample is calculated by adding $x_c = x_a + x_b$. The SMOTE method allows you to specify the number of new samples to be artificially generated. To ensure the quality control of training models, the following model construction scheme is proposed. A training set and a test set are preliminarily created from the entire data set. The model is trained using a cross-validation procedure, in which K blocks (folds) are divided into the data set. The learning process is performed K times on different training sets consisting of K-1 blocks. To improve the balance of the training sample classes, they are extended by artificially generating samples that are "similar" to those in the bankruptcy class, but do not duplicate them. For generation, the SMOTE method is used. To control the training quality, a remaining The validation set is used to test the quality of the training. The test set is then used to evaluate how well-trained the model is. As part of cross-validation, K models were created and averaged to gain a final prediction.. Since, as part of the cross-validation procedure, K models were built, we obtain K predictions on the test set, which are then averaged.

Because the SMOTE method is a technique that has built-in advantages such as being organically integrated, it holds many benefits and therefore, is included in the training scheme model as the extreme gradient boosting method. This is a naturally advantageous version of extreme gradient boosting which, in turn, leads to increased accuracy in prediction models.

2.3. The value of quality

One common way to evaluate the quality of a classification model is accuracy, which is the number of objects that are accurately classified. However, this measurement isn't very convenient when the classes involved have uneven distributions; in such cases, we use precision and recall separately and calculate the proportion of positive calls that are accurate and accurate out of all those in the positive group. Precision can be interpreted as the number of correct calls (true positives) in a class that have been identified by the classifier, and recall identifies how many there are total among those within a class that are known to be true positives. The precision and recall value is determined by the class ratio; on the other hand, precision does not depend on recall. For this reason, they are still effective when there is an unbalanced set of conditions.

When designing a machine learning, you can use a metric as a guide for how to set its hyperparameters. The metric improvement that we expect with the test set reflects the performance of the full model. As such, in the metric it is proposed to use the F1 score metric, which represents the harmonic mean precision and recall. To calculate the values for precision, recall, F1 and accuracy, you have to choose a threshold at which the decision becomes 0 or 1. A natural threshold is 0.5 but it isn't always optimal, especially when there's an unbalanced class distribution. One way to evaluate the model as a whole without being tied to a specific threshold is the AUC metric, which is numerically equivalent to the area under the accuracy curve.

3. Experimental study

3.1. Data set

In this article, the data was taken from the UCI1 machine learning repository. These were originally extracted from Emerging Markets Information Services2, which is a database containing information on emerging markets around the world. The study looked at bankruptcy filings during the years 2000–2012. For still-operating companies, data was collected from 2007 to 2013.

In total, five datasets are provided containing financial indicators of some enterprises and a label indicating the bankruptcy status after 1, 2, 3, 4 and 5 years. The metadata of the datasets are presented in the table.

Table 1. Key features of the data set

Characteristic	Bankruptcy label				
	1 year	2 year	3 year	4 year	5 year
Number of enterprises	5810	9892	10513	10163	7037
Share of bankrupt enterprises	6,93	5,26	4,71	3,93	3,85
Proportion of missing values	1,21	1,37	1,44	1,83	1,27

Thus, we have 5 sets that are not equivalent. Three datasets labeled with bankruptcy status at 2, 3, and 4 years contain a similar number of businesses around 10,000. The 1-year bankruptcy-labeled dataset has 40% fewer companies, and the 5-year bankruptcy-tagged dataset includes 30% fewer companies than the 2-year bankruptcy-tagged datasets, 3 and 4 years. In the sets, each enterprise is described by 64 financial indicators that measure profitability, the enterprises' assets and their cash flows. A full description of the indicators can be found in [15], which examines the effectiveness of machine learning methods in building bankruptcy prediction models. However, the studies conducted in this paper have a number of shortcomings. First, the entire set is used for the cross-validation procedure. This rating depends only on the quality of the model training, especially when using methods such as Random Forest and Extreme Gradient.

We'll be taking a step back with our model, in order to see how well it generalizes to new and previously unseen data points. To do this, we're going to employ a test set that won't be trained on. Second, the AUC metric is used as a quality measure, which is not sufficient for a complete assessment of the quality of enterprise bankruptcy prediction. So we will use AUC to use the F1 metric for each class.

3.2. Data preprocessing

When building the models, the data sets described above are used, which are divided as follows. The whole set is divided into two parts: training (70%) and test (30%). Allocation of a sufficiently large test set makes it possible to increase the reliability of the assessment of the generalizability of the model.

One of the first tasks in data preprocessing is filling in missing values. There are several strategies to fill in the gaps. The most popular strategies are strategies based on assigning missing values mean (mean) or median (median) values. When using decision trees, a missing value flagging strategy is effective, such as assigning large negative numbers (constants) to the missing values.

Table 2 shows the results of comparing different missing value filling strategies for the dataset labeled bankruptcy status 5 years. Quality measurement for models obtained by the XGB method. Columns in the table titled "N" contain F1 and AUC values for non-bankrupt firms and columns labeled "B" are F1 and AUC values for bankrupt firms.

Table 2 Comparing decision quality for different strategies

Strategie	F1		AUC	F1		AUC
	N	B		N	B	
Mean	0,9801	0,6695	0,953	0,9877	0,6525	0,972
Median	0,9790	0,6167	0,947	0,9873	0,6318	0,974
constant	0,9807	0,6755	0,951	0,9914	0,7163	0,979

Analysis of the results presented in the tables shows that the constant strategy, which fills in the missing values with constant numbers equal to -99999, is the most effective. Research on other data sets also shows a small advantage of the constant strategy. Research done on other data sets confirms the preference for the constant strategy over other strategies.

The second task, which is solved in the stage of data preparation, is to select the most informative features from the total number. The presence of some non-informative features in the data leads to the retraining of the model and a decrease in the accuracy of the results obtained.

In the table. 3 shows a quality comparison between the models obtained with a different number of features that were sorted by the significance coefficient

Table 3 Comparison of the quality of solutions for various

Număr semne	F1		AUC	F1		AUC
	N	B		N	B	
64	0,990	0,685	0,941	0,99	0,715	0,959
50	0,990	0,692	0,949	0,99	0,695	0,960
40	0,990	0,694	0,952	0,99	0,733	0,960
30	0,990	0,694	0,945	0,99	0,723	0,941

The top 10 most important signs include the following signs:

- 1) the ratio between profit from exploitation activities and financial expenses;
- 2) current liquidity rate;
- 3) the ratio between operating expenses and total debts;
- 4) logarithm of total assets;
- 5) the relationship between the company's current assets and its total liabilities, which includes cash, short-term investments and accounts receivables minus short-term liabilities;
- 6) the ratio between total costs and total sales;
- 7) the ratio between sales revenue and total assets;
- 8) the ratio of receivables multiplied by 360 to sales revenue;

- 9) the ratio between the difference between current assets and reserves and long-term liabilities;
 10) the ratio between the amount of gross profit, other assets and transactions, financial costs to total assets.

These ratios measure not only the profitability of businesses, but also the assets of businesses and their cash flows. Analyzing the results of the conducted studies, we can say that these indicators appear in the list of important indicators for each data set. In addition to these 10 indicators, 35 more indicators that positively affect the quality of models were selected.

3.3. Comparing the quality of models built by different methods

You can build models using the following methods: linear discriminant analysis, logistic regression, random forest, and extreme gradient boosting. For even better accuracy, try extreme gradient boosting and SMOTE. When using the LR, RF, XGB, XGB-Sm methods, the search for the best hyper-parameters was performed by listing their values in the given ranges. During the study, 10 experiments were performed. During the experiment, each set is divided into training (70%) and test (30%) sets. When building a model on a training set, a 10-block cross-validation procedure is used. The quality of the short term forecast of bankruptcy's probability should be estimated on a validation set. As a rule, the estimation happens in such a way that the model F is inverted proportional to F1 and P. Quality values obtained over 10 experiments are averaged. In the table. Figures 4–8 show the model quality metrics obtained by various methods: the average value of the F1 metric for operating (N) and bankrupt (B) businesses, and the average AUC value and precision metrics.

Table 4 Comparison of model quality with a 1-year forecast horizon

METHOD	Validation set				test set			
	F1		AUC	ACC	F1		AUC	ACC
	N	B			N	B		
LDA	0,9562	0,5075	0,855	0,9467	0,9566	0,507	0,85	0,937
LR	0,9613	0,4751	0,709	0,9556	0,9609	0,319	0,71	0,944
RF	0,9538	0,5667	0,919	0,9433	0,9569	0,619	0,93	0,939
XGB	0,9713	0,6871	0,953	0,9547	0,9730	0,713	0,96	0,957
XGB-Sm	0,9714	0,7152	0,957	0,9552	0,9729	0,753	0,96	0,969

In the table 4 shows the results obtained on the dataset labeled with bankruptcy status after 1 year. The share of bankrupt enterprises in the validation and test sets are 0.0696 and 0.0688, respectively. Thus, if the decision to use a simple solution (all enterprises are active), then the values of the precision metric will be 0.9304 and 0.9312, respectively. We assume that these values are thresholds, i.e. forecasts, whose accuracy measure is below the threshold values, will be considered bad.

The quality of the models obtained by the LDA and RF methods slightly exceeds the threshold values, while the AUC metric has quite high values. The use of the logistic regression method makes it possible to obtain forecasts whose precision metric exceeds the thresholds by almost 1.5%, and the AUC value is much below this in the case of applying the LDA and RF methods.

Logistic regression models accurately classify 106 of the 288 bankrupt firms in the validation set and 39 of the 122 in the test set. The 3806 successful plants and 1636 successful plants are also correctly classified by logistic regression models for the validation and test sets, respectively. Of the enterprises operating the validation and test sets, 3750 out of 3849 and 1605 out of 1651 are correctly classified, respectively.

The models built with RF allow correct classification of 188 out of 288 and 88 out of 122 validation and test set enterprises, respectively. From the enterprises operating the validation and testing sets correctly 3673 out of 3849 and 1577 out of 1651 respectively are classified.

The models obtained with XGB and XGB Sm demonstrate the highest accuracy. Improving the balance of the training datasets improves the forecast quality, such that the F1 quality metric for bankrupt firms increases by 2.8% and 3.9% for the validation and testing datasets, respectively. The F1 metric for operating businesses does not change. The optimal balance ratio is 0.09.

Table 5 Comparison of model quality with a 2-year forecast horizon

METHOD	Validation set				Test set			
	F1		ACU	ACCURACY	F1		AUC	ACURATETE
	N	B			N	B		
LDA	0,9763	0,5038	0,815	0,9629	0,9742	0,4536	0,825	0,9479
LR	0,9752	0,4329	0,680	0,9607	0,9753	0,3787	0,658	0,9516
RF	0,9831	0,6274	0,893	0,9710	0,9829	0,5897	0,886	0,9642
XGB	0,9849	0,6201	0,944	0,9791	0,9845	0,5913	0,934	0,9673
XGB-Sm	0,9845	0,6861	0,952	0,9810	0,9833	0,6364	0,944	0,9684

Table 5 shows the quality metrics of the models obtained on sets labeled with the bankruptcy status of the enterprises after 2 years. The accuracy metric thresholds for the validation and test sets are 94.61% and 95.03%, respectively.

As the table below shows, on the test set, the LDA models have a low accuracy score and the LR models have a slight increase in accuracy. The models built by the XGB and XGB-Sm methods show the highest values of the quality metrics. Improving the balance of the training datasets improves the forecast quality, for example, the F1 quality measure for bankrupt firms increases by 6.6% and 4.7% for the validation and test sets, respectively.

Although the F1 metric for operating businesses is slightly reduced (0.04% and 0.01%). The optimal balance coefficient of the class is 0.09

Table 6 Comparison of model quality with a 3-year forecast horizon

METHOD	Validation set				Test set			
	F1		AUC	ACCURACY	F1		AUC	ACURATETE
	N	B			N	B		
LDA	0,9426	0,3421	0,780	0,9116	0,9443	0,3182	0,780	0,9133
LR	0,9577	0,2720	0,651	0,9381	0,9550	0,2327	0,653	0,9320
RF	0,9695	0,5207	0,875	0,9607	0,9652	0,4186	0,846	0,9514
XGB	0,9742	0,5591	0,921	0,9695	0,9728	0,4274	0,902	0,9657
XGB-Sm	0,9752	0,6534	0,928	0,9716	0,9723	0,5597	0,909	0,9660

Table 6 shows the quality metrics of the models obtained on sets labeled with the bankruptcy status of the enterprises after 3 years. The accuracy metric thresholds for the validation and test sets are 95.25% and 95.36%, respectively.

As can be seen from the table, the accuracy metric of the models built by the LDA and LR methods is the test set and is below the threshold values. The models built using the RF method, when classifying enterprises in the test set, show results where the proportion of correctly classified enterprises is below the threshold of average values, which also indicates the low quality of the model.

The models built by the XGB and XGB-Sm methods show the highest quality metric values. By increasing the balance of the datasets and adjusting the F1 quality metric, our forecast quality improves, albeit at an uneven rate. For bankrupt firms, the F1 metric increases by 9.5% and 13.1% on our validation and test sets respectively, but operational enterprises actually see a 1% decrease on our validation set and a 0.05% decrease on our test set. The best balance coefficient for this dataset is 0.06.

Table 7 Comparison of model quality with a 4-year forecast horizon

METHOD	Validation set				Test set			
	F1		AUC	ACCURACY	F1		AUC	ACURATETE
	N	B			N	B		
LDA	0,9563	0,3389	0,721	0,9458	0,9793	0,2800	0,721	0,9410
LR	0,9641	0,2552	0,601	0,9599	0,9909	0,2083	0,606	0,9626
RF	0,9722	0,5343	0,865	0,9756	0,9955	0,5275	0,850	0,9718
XGB	0,9768	0,5767	0,914	0,9844	0,9982	0,5333	0,937	0,9771
XGB-Sm	0,9774	0,6647	0,921	0,9957	0,9982	0,6316	0,939	0,9780

Table 7 presents the quality metrics of the models obtained on sets labeled with the bankruptcy status of the enterprises after 4 years. The accuracy metric threshold values for the validation and test datasets are 95.87% and 96.53%, respectively. As can be seen from the table, the metric precision of the quality of the models built by the LDA and LR methods, on the validation and test sets, have a value below the threshold values.

The XGB and XGB-Sm models produce higher quality scores than those of the validation and test sets. Improving the balance of the dataset is necessary to improve forecast quality, which resulted in a F1 quality metric that increased 8.8% on the validation set, and 9.8% on the test set (compared to no change for existing businesses). A class balance ratio of 0.06 led to optimal results.

Table 8 Comparison of model quality with a 5-year forecast horizon

METHOD	Validation set				Test set			
	F1		ACU	ACCURACY	F1		AUC	ACURATETE
	N	B			N	B		
LDA	0,9675	0,4828	0,842	0,9469	0,9812	0,5870	0,924	0,9540
LR	0,9714	0,5194	0,882	0,9542	0,9820	0,5629	0,893	0,9754
RF	0,9715	0,4599	0,908	0,9542	0,9808	0,5185	0,942	0,9730
XGB	0,9809	0,6826	0,937	0,9723	0,9809	0,6947	0,942	0,9706
XGB-Sm	0,9810	0,7175	0,943	0,9727	0,9914	0,7552	0,969	0,9844

Table 8 shows the quality metrics of the models obtained on sets labeled with the bankruptcy status of enterprises after 5 years. The following table illustrates the accuracy of the LR method under different threshold values:

The threshold for the precision metric is 96.28%, while that of validation and test sets is 95.83%. As can be seen from the table, when tested on the validation set, LR has an accuracy which falls below these thresholds. However, this is true only for the first tests done. The models built by the XGB and XGB-Sm methods show the highest quality metric values. Improving the balance of the training datasets improves the forecast quality, but only slightly. The optimal balance ratio in classes is 0.07.

Thus, the models built by the LDA and LR methods provide the lowest prediction accuracy, which is often below the threshold values. The models built by the RF method, although they give a higher prediction, in some cases do not allow obtaining good quality models. The models built by XGB-Sm demonstrate the highest quality metric values in both validation and test sets. It should be noted that if the number of artificial samples, generated by the SMOTE method exceeds the initial number of samples, the quality of the models, built by the extreme increase of the gradient, decreases significantly.

Conclusion

The article considers the problem of predicting bankruptcy based on financial factors. To solve the mentioned classification problem, we use a model built with the help of extreme gradient boosting. Increasing the extreme gradient resulted in a better result than linear discrimination, logistic regression - well-known methods of financial forecasting. The paper proposes an extremal gradient extension that improves training set class imbalance using the SMOTE method. Applying this approach led to an improvement in forecast quality.

Future studies are planned to be conducted in the direction of considering growth rates of financial indicators when building models to assess the effect of time on the probability of bankruptcy of enterprises.

Acknowledgements

This work is supported by project POCU/380/6/13/125031 „DECIDE - Development through entrepreneurial education and innovative doctoral and postdoctoral research", co-financed by the European Social Fund under the Human Capital Operational Program 2014-2020.

References

1. Altman E.I. *Financial Ratios, Discriminant Analysis and the Prediction of Corporate Bankruptcy* // *The Journal of Finance*. 1968. Vol. 23, no. 4. P. 589–609. DOI: 10.1111/j.1540-6261.1968.tb00843.x.
2. Lugovskaya L. *Predicting Default of Russian Smes on the Basis of Financial and Nonfinancial Variables* // *Journal of Financial Services Marketing*. 2010. Vol. 14, no. 4. P. 301–313. DOI: 10.1057/jfsm.2009.28.
3. Deakin E. *A Discriminant Analysis of Predictors of Business Failure* // *Journal of Accounting Research*. 1972. Vol. 10, no. 1. P. 167–179. DOI: 10.2307/2490225.
4. Antunesa F., Ribeiroa B., Pereirab F. *Probabilistic Modeling and Visualization for Bankruptcy Prediction* // *Applied Soft Computing*. 2017. Vol. 60. P. 831–843. DOI: 10.1016/j.asoc.2017.06.043.
5. Ohlson J.A. *Financial Ratios and the Probabilistic Prediction of Bankruptcy* // *Journal Of Accounting Research*. 1980. Vol. 18, no. 1. P. 109–131. DOI: 10.2307/2490395.
6. Martin D. *Early Warning of Bank Failure: A Logit Regression Approach* // *Journal of Banking & Finance*. 1977. Vol. 1, no. 3. P. 249–276.
7. Wiginton J.C. *A Note on the Comparison of Logit and Discriminant Models of Consumer Credit Behavior* // *Journal of Financial and Quantitative Analysis*. 1980. Vol. 15. P. 757–770. DOI: 10.2307/2330408.
8. Begley J., Ming J., Watts S. *Bankruptcy Classification Errors in the 1980s: an Empirical Analysis of Altman's and Ohlson's Models* // *Review of Accounting Studies*. 1996. Vol. 1, no. 4. P. 267–284. DOI: 10.1007/bf00570833
9. Wilson R.L., Sharda R. *Bankruptcy Prediction Using Neural Networks* // *Decision Support Systems*. 1994. Vol. 11, no. 5. P. 545–557. DOI: 10.1016/0167-9236(94)90024-8.
10. Tam K.Y., Kiang M.Y. *Managerial Applications of Neural Networks: the Case of Bank Failure Predictions* // *Management Science*. 1992. Vol. 38, no. 7. P. 926–947. DOI: 10.1287/mnsc.38.7.926.
11. Altman E.I., Marco G., Varetto F. *Corporate Distress Diagnosis: Comparisons Using Linear Discriminant Analysis and Neural Networks (the Italian Experience)* // *Journal of Banking & Finance*. 1994. Vol. 18, no. 3. P. 505–529. DOI: 10.1016/0378-4266(94)90007-8.
12. Ciampi F., Vallini C., Gordini N. *Using Artificial Neural Networks Analysis for Small Enterprise Default Prediction Modeling: Statistical Evidence from Italian Firms* // *Oxford Business & Economics Conference Proceedings, Association for Business and Economics Research, ABER*. 2009. Vol. 1. P. 126.
13. Wei L., Li J., Chen Z. *Credit Risk Evaluation Using Support Vector Machine with Mixture of Kernel* // *Proceedings of the 7th International Conference on Computational Science. Lecture Notes in Computational Science and Engineering*. 2007. Vol. 4488. P. 431–438. DOI: 10.1007/978-3-540-72586-2_62.
14. Härdle W.K., Lee Y.J., Schäfer D. *The Default Risk of Firms Examined with Smooth Support Vector Machines* // *Discussion Papers, German Institute for Economic Research*. 2007. no. 757. P. 1–30. DOI: 10.2139/ssrn.2894311.
15. Zieba M., Tomczak S.K., Tomczak J.M. *Ensemble Boosted Trees with Synthetic Features Generation in Application to Bankruptcy Prediction* // *Expert Systems with Applications*. 2016. Vol. 58. P. 93–101. DOI: 10.1016/j.eswa.2016.04.001.
16. Xia Y., Liu C., Li Y., Liu N. *A Boosted Decision Tree Approach Using Bayesian HyperParameter Optimization for Credit Scoring* // *Expert Systems With Applications*. 2017. Vol. 78. P. 225–241. DOI: doi.org/10.1016/j.eswa.2017.02.017.
17. Zhou L. *Performance of Corporate Bankruptcy Prediction Models on Imbalanced Dataset: The Effect of Sampling Methods* // *Knowledge-Based Systems*. 2013. Vol. 41. P. 16–25. DOI: 10.1016/j.knsys.2012.12.007.

18. Kim T., Ahn H. A Hybrid Under-Sampling Approach for Better Bankruptcy Prediction // *Journal of Intelligent Information Systems*. 2015. Vol. 21, no. 2. P. 173–190. DOI: doi.org/10.13088/jiis.2015.21.2.173.
19. Véganzones D., Séverina E. An Investigation of Bankruptcy Prediction in Imbalanced Datasets // *Decision Support Systems*. 2018. no. 112. P. 111–124. DOI: [10.1016/j.dss.2018.06.011](https://doi.org/10.1016/j.dss.2018.06.011).
20. Chen T., Guestrin C. XGBoost: A Scalable Tree Boosting System // *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*. ACM. 2016. P. 785–794. DOI: [10.1145/2939672.2939785](https://doi.org/10.1145/2939672.2939785).
21. Friedman J.H. Stochastic Gradient Boosting // *Computational Statistics and Data Analysis*. 2002. no. 38. P. 367–378. DOI: [10.1016/j.eswa.2016.04.001](https://doi.org/10.1016/j.eswa.2016.04.001)