



Credit Risk Detection in Peer-to-Peer Lending Using CatBoost

Fadhlurrahman Akbar Nasution¹, Siti Saadah², Prasti Eko Yunanto³

^{1,2,3}Department of Informatics, Informatics, Telkom University, Bandung, Indonesia

¹fadhlurrahmanakbar@student.telkomuniversity.ac.id, ²sitisaadah@telkomuniversity.ac.id, ³gppras@telkomuniversity.ac.id

Abstract

P2P (Peer-to-peer) lending has gained popularity among private borrowers, small businesses, and MSMEs due to its ability to provide direct access to loans without the strict requirements imposed by traditional banks and financial institutions. However, P2P lending faces a significant challenge in terms of credit risk, resulting in a high rate of loan repayment failures. To address this issue, the study aimed to develop a credit risk detection system using a loan dataset obtained from the Bondora company by implementing one of the gradient boosting algorithms which are called the CatBoost (Categorical Boosting) method. The performance of the CatBoost algorithm was evaluated using ROC (Receiver Operating Characteristics) curves and AUC (Area Under Curve). Three scenarios were considered, and the results revealed that scenario 2, with a data splitting ratio of 90:10, achieved the best outcome with an AUC value of 0.80329. This outperformed scenario 1, with a data splitting ratio of 80:20 and an AUC value of approximately 0.789583, as well as scenario 3, with a data splitting ratio of 70:30 and an AUC value of around 0.781066.

Keywords: P2P lending; credit risk detection; catboost; AUC; ROC

1. Introduction

P2P lending has become increasingly popular among individuals, small businesses, and MSMEs because it offers more flexible requirements and less strict criteria in contrast to traditional banks and financial institutions [1]. However, P2P lending has a high failure rate for borrowers to repay their loans and P2P lending also has low-interest rates, which is known as credit risk [1]–[3]. Credit risk is a major problem in P2P Lending because payers who fail to make loan payments make the lender suffer losses [2]. To minimize credit risk, a system is needed that can assist in determining credit risk.

To overcome this challenge, a paper survey by M. Bazarbash [4] suggests utilizing machine learning to assess credit risk in P2P lending. Machine learning involves the application of sophisticated algorithms that can be executed by machines to identify patterns in data, with the primary aim of making predictions. Machine learning models are specifically designed to analyze extensive information from diverse data sources, enabling them to detect patterns that conventional econometric models may overlook. By comprehending the complex and nonlinear relationships between risk factors and credit risk outcomes, machine learning techniques offer a promising solution to address credit risk in P2P lending. They have the capability to mitigate

the issue of information asymmetry, which is at the core of challenges in P2P lending. Consequently, the detection of credit risk in P2P lending can play a crucial role in minimizing loan repayment failures.

Many studies about credit risk prediction using machine learning approaches have been observed. In 2018, J. Mezei [5] conducted a study "Credit risk evaluation in peer-to-peer lending with linguistic data transformation and supervised learning". The study employed fuzzy sets-based linguistic data transformation combined with a neural network method using the Bondora dataset spanning from February 2, 2009, to June 17, 2017. The study achieved an AUC value of 0.855. In 2019, Linnea Machado [6] conducted a study "Credit risk modeling and prediction: Logistic regression versus machine learning boosting algorithms". The research utilized two datasets and employed various algorithms including CatBoost, Logistic Regression, and XGBoost. The findings indicated that the CatBoost algorithm outperformed the others. In dataset 1, CatBoost achieved an AUC value of 0.863 and an accuracy of 0.836 on the training data, while on the test data, it achieved an AUC value of 0.731 and an accuracy of 0.817. When applied to dataset 2, CatBoost achieved an AUC value of 1 and an accuracy of 0.999 on the training data, and an AUC of 0.925 and an accuracy of 0.946 on the test data.

In the following year, in 2022, Nghia Nguyen [7] conducted a study titled "A Proposed Model for Card Fraud Detection Based on CatBoost and Deep Neural Network." This study presented two models: a neural network model and a CatBoost model. The experiment utilized a dataset from IEEE CIS, which consisted of real-world transactions provided by Vesta Corporation. The results demonstrated that the CatBoost model outperformed the neural network model, achieving an AUC value of 0.974, whereas the neural network model achieved an AUC value of 0.84. During the same year, Xingyun Li [8] conducted a study titled "Prediction of Loan Default Based on Multi-Model Fusion." This research introduced a novel fusion model that combined Logistic Regression, Random Forest, and CatBoost techniques to predict loan defaults. The results indicated that the multi-model fusion, which integrated multiple techniques, slightly outperformed the CatBoost model. The CatBoost model achieved an AUC value of 0.992, whereas the multi-model fusion proposed by the study achieved an AUC value of 0.994.

Based on insights from prior research, the present study adopts the CatBoost model due to its strong performance. The evaluation metrics employed will be AUC (Area Under Curve) and ROC (Receiver Operating Characteristic) curve, as these were utilized in previous studies to assess the effectiveness of credit risk detection. The dataset utilized in this study originates from the Bondora company. The primary objective of this research is to identify credit risk in the context of P2P lending.

2. Research Methods

In this study, the objective is to develop a credit risk assessment system in P2P lending that harnesses the CatBoost algorithm, the whole process presents in Figure 1.

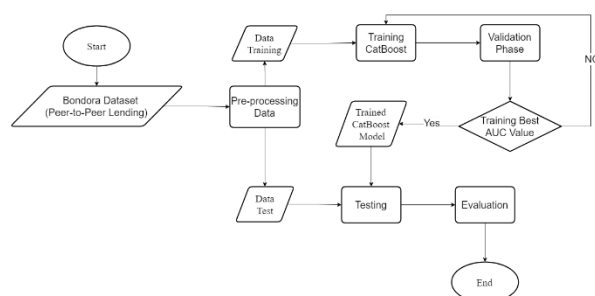


Figure 1. System Design Flowchart

Based on the system design in Figure 1, the first step involves preprocessing the P2P Lending dataset. Subsequently, the dataset is divided into two subsets: training data and test data, which are used for the training and testing stages, respectively. Following the training phase, a validation step employing three-fold

cross-validation is conducted to determine the optimal AUC result. The model yielding the best AUC result is then utilized in the testing phase. Then, the performance will be evaluated in the testing phase using the ROC curves and AUC.

2.1 Credit Risk in P2P Lending

P2P lending is an internet-based loan service that establishes a direct connection between borrowers and lenders [9]. The emergence of P2P lending has eliminated the necessity for traditional banks to serve as intermediaries for borrowers [10]. Small businesses previously encountered challenges in accessing loans from conventional banks due to factors such as elevated default rates, limited data accessibility, and the perception that small business loans were not lucrative [10]. The accessibility and low-interest rates offered by P2P lending increase the credit risk, as the large number of borrowers raises concerns about loan repayment defaults [11], [12]. Furthermore, the inadequate implementation of risk control measures by P2P lending platforms has exacerbated the rise in credit risk [13].

Credit risk poses a significant challenge in the realm of P2P lending as borrowers frequently default on loan payments, leading to financial losses for lenders [14]–[16]. The problem of credit risk is compounded by the information imbalance between borrowers and lenders, particularly in the context of P2P lending, where unsecured loans are commonly involved [17]. Consequently, there is a pressing need for an effective and precise credit risk assessment method that leverages machine learning techniques. Such an approach would serve to mitigate credit risk and ensure the sustainable growth of the P2P lending sector [1].

2.2 Bondora Dataset (Peer-to-Peer Lending Dataset)

The dataset utilized is a publicly available loan dataset from the Bondora company¹. The dataset is publicly available and always updated every day [5]. The loan dataset contained the records of transactions from February 28, 2009, to November 14, 2022, in euros (EUR). The dataset contained 263,213 records and 112 column attributes, so that column attributes will be selected according to the prediction of failure in loan payments. The selected attribute columns in Table 2 are based on the analysis of variables that will be utilized for identifying loan defaults.

The characteristics of the dataset are categorical and numerical data, with categorical data located in the New Credit Customer, Country, Status, Marital Status, Use of Loan, and Employment Status columns. Meanwhile, numerical data is in the Age, Applied Amount, Interest, Loan Duration, and Income Total (see Table 1).

¹ <https://www.bondora.com/en/public-reports#shared-legend>

Table 1. Sample Data

New Credit Customer	Age	Country	Applied Amount	Interest	Loan Duration	Use Of loan	Marital Status	Employment Status	Income Total	Status
True	29	EE	1805	29.62	12	3	1	3	446	Repaid
True	39	FI	10630	53.01	48	-1	-1	-1	2400	Late
True	38	ES	10630	57.13	36	7	3	6	2330	Repaid
True	47	EE	2975	32.18	60	-1	-1	-1	1050	Late
True	40	EE	2000	36.00	48	0	2	5	913	Repaid

Table 2. Attribute Description Columns Used

Feature Column	Description
New Credit Customer	Customers who possess a prior credit record with Bondora
Age	Borrower age
Country	Borrower's country of origin
Applied Amount	The borrower submitted the loan amount at the start
Interest	The maximum acceptable interest rate
Loan Duration	The duration of the loan in months
Use Of Loan	The purpose of the loan
Marital Status	Marital status categories (1 for Married, 2 for Living at Home, 3 for Single, 4 for Divorced, 5 for Widows)
Employment Status	Occupational status categories (1 for Unemployed, 2 for Part Time, 3 for Full Time, 4 for Self-employed, 5 for Entrepreneur, 6 for Retired)
Income Total	The total salary of the borrower
Status	Current loan status

2.3 Pre-processing

In the first stage, pre-processing will be conducted before the data is used to build the model. The pre-processing steps are outlined in Figure 2.

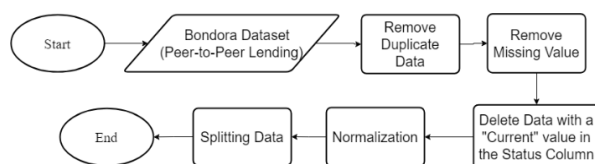


Figure 2. Preprocessing Flowchart

Initially, given the substantial amount of data, the elimination of duplicate entries and missing values will be removed, as this is a widely employed technique to enhance the model's optimization at a later stage. Subsequently, the analysis will specifically focus on data labeled as "repaid" and "late", with the data labeled as "current" being excluded from the analysis. Next, data normalization will be applied to transform categorical data into numerical values, while numerical data will be scaled using a min-max scaler. Finally, the dataset will be split into training and test sets in a specific ratio to ensure optimal outcomes.

2.4. CatBoost

Binary decision trees are utilized as the fundamental predictors in CatBoost, which is an implementation of gradient boosting [18]. The challenges arising from diverse features, noisy data, and complex relationships are effectively addressed by gradient boosting, which is a robust machine learning technique [19]. Notably, CatBoost algorithm excels not only in handling numeric features but also in effectively managing categorical features [20]. It employs an unconstrained decision tree structure where all non-terminal nodes at the same level of the tree possess identical splitting criteria. This ensures that the length of the path from the root node to each leaf is equal to the depth of the tree [21]. An illustration of the CatBoost can be seen in Figure 3.

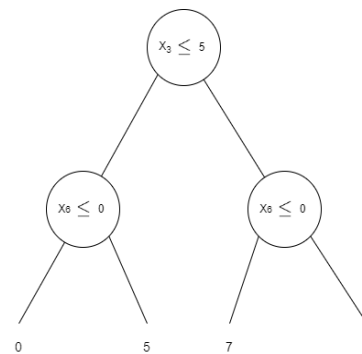


Figure 3. Illustration CatBoost[21]

CatBoost employs ordered target statistics to handle categorical features [18]. This approach involves estimating the target value for each category by calculating the corresponding output value. The specific mathematical equation for this process is provided in Equation **Error! Reference source not found.** [6].

$$E(Y_i | x_i = x_{i,k}) \quad (1)$$

The average of the output Y_i is denoted by $x_{i,k}$, that is value of the i -th categorical variable in the k -th exercise observations from the given dataset. The equation can be seen in Equation **Error! Reference source not found.** [6].

$$\hat{x}_{i,k} = \frac{\sum_{j=1}^n [x_{i,j} = x_{i,k}] \cdot y_j + aP}{\sum_{j=1}^n [x_{i,j} = x_{i,k}] + a} \quad (2)$$

If the dot $[\cdot]$ corresponds to the inversion brackets, then the value $[x_{i,j} = x_{i,k}]$ equals to 1 and the value $x_{i,j}$ is equal to $x_{i,k}$, otherwise the value is equal to 0.

Parameter P is the prior probability to determine the default class using parameter $a > 0$ [18].

2.5. Training

Once the data has been pre-processed, the training phase for the CatBoost model commences. The process begins by determining the value of N, representing the number of trees trained in the CatBoost model. Subsequently, a random permutation is performed, followed by the creation of a new tree that predicts the permutations based on the trained model. The permutation estimates for each model are recalculated, and the model with the best permutation will be used for the testing stage. The training process will be conducted until N iteration. A visual representation of the CatBoost training stages can be found in Figure 4.

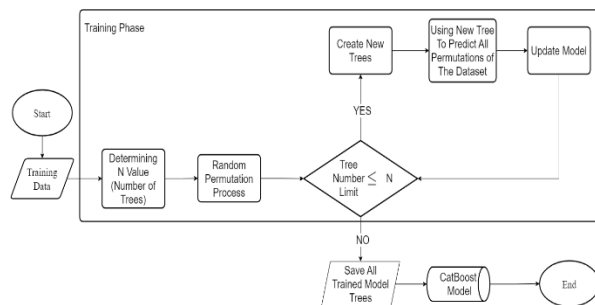


Figure 4. CatBoost Training Flowchart

2.6. Testing

Once the CatBoost training is completed, the model will proceed to the testing phase. The steps involved in CatBoost testing are illustrated in Figure 5.

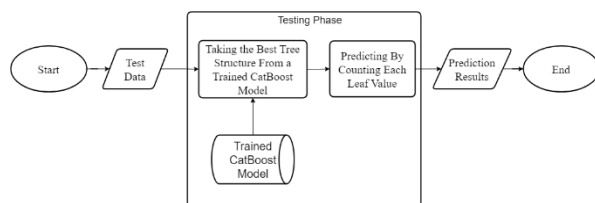


Figure 5. CatBoost Testing Flowchart

During the model testing phase, the parameters obtained from the best tree structure during training are

Table 3. Sample Data Preprocessed

New Credit Customer	Age	Country	Applied Amount	Interest	Loan Duration	Use Of loan	Marital Status	Employment Status	IncomeTotal	Status
0	0.368421	0	0.378740	0.035781	0.495798	0	0	0	0.001078	0
1	0.315789	0	0.031355	0.081283	0.092437	0	0	0	0.000692	0
0	0.500000	0	0.031292	0.063411	0.294118	0	0	0	0.000695	0
1	0.526316	1	0.265018	0.125254	0.294118	0	0	0	0.002765	1
1	0.657895	1	0.198105	0.217098	0.495798	0	0	0	0.016304	0

3. Results and Discussions

To fulfill the study's objectives, three scenarios were executed to assess: 1. The effectiveness of the CatBoost algorithm's hyperparameters in credit risk classification; 2. A comparison of different training and

employed by the CatBoost model. By traversing the tree and counting the number of leaves, predictions are made. The prediction outcomes are then evaluated using AUC and ROC curves to evaluate the effectiveness of CatBoost in credit risk assessment for P2P lending.

2.7. Model Evaluation

The ROC curve is a visual representation utilized to evaluate the classifier's performance [6], [22]. In the ROC curve, The Sensitivity or True Positive Rate (TPR), which represents the proportion of correct predictions classified as true, is depicted on the Y-axis [22]. On the other hand, the False Positive Rate (FPR), which represents to the ratio of incorrect predictions classified as correct predictions, is depicted on the X-axis. The ROC curve is utilized to visually represent the CatBoost model's capacity to precisely classify credit risk by evaluating the ratio of accurate and inaccurate predictions.

$$TPR = \frac{TP}{TP + FN} \quad (1)$$

$$FPR = \frac{FP}{FP + TN} \quad (2)$$

Then there is an important index of ROC, namely AUC, which is the value of the area between the ROC curve and the abscissa [22]. A better credit risk score is indicated by a higher value of the AUC (Area Under the Curve), which ranges from 0 to 1 [22].

TPR and FPR can be calculated using Equation **Error! Reference source not found.** and **Error! Reference source not found.** In these equations, the count of positive instances that are accurately classified is denoted by TP (True Positive), The count of positive instances that are mistakenly classified as negative is denoted by FP (False Positive), The count of negative instances that are accurately classified is denoted by TN (True Negative), and the count of negative instances that are inaccurately classified is denoted by FN (False Negative).

testing data ratios; 3. The system's performance measurement using AUC and ROC metrics.

3.1 Preprocessed Dataset

After the pre-processing stage, the dataset is reduced to 171,257 instances, where the data consist of 94,258

rows Repaid and 76,999 rows Late status. The data distribution can be seen in Figure 6.

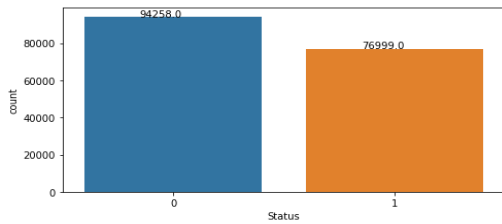


Figure 6. Label Distribution

The status values are encoded, where "Repaid" is represented as "0" and "Late" is represented as "1". In the process of normalizing categorical attributes, "New Credit Customer" and "Country" are manually encoded, while "Marital Status", "Use of Loan", and "Employment Status" attributes are encoded using the sklearn library. The dataset will be divided into three scenarios: 70:30, 80:20, and 90:10. These scenarios were chosen to assess the CatBoost model's ability to identify patterns in the training data and produce good results when tested on the test data. The preprocessed data is presented in Table 3.

3.2 Result

The experiment of this study was conducted to obtain the best parameters by utilizing the grid search method as shown in Table 4. The grid search method is utilized to search for optimal results for the depth and learning rate parameters. The iteration parameter is manually set to 500 for each scenario. Additional parameters, such as the random_seed parameter and cat_features, are included to enhance the performance of CatBoost.

Table 4. Experiment Scenario

Parameter	Value		
	Scenario 1 (80:20)	Scenario 2 (90:10)	Scenario 3 (70:30)
iteration	500	500	500
loss_function	Logloss	Logloss	Logloss
random_seed	1	1	1
verbose	100	100	100
eval_metric	AUC	AUC	AUC
task_type	CPU	CPU	CPU
early_stopping_rounds	100	100	100
cat_features	[0, 1, 2, 3, 4, 5, 6, 7, 8, 9]	[0, 1, 2, 3, 4, 5, 6, 7, 8, 9]	[0, 1, 2, 3, 4, 5, 6, 7, 8, 9]
depth	8	8	7
learning_rate	0,1	0,1	0,1

In this study, the random seed parameter is set to 1 for each scenario, while the cat_features parameter is determined based on the total number of attribute columns in the dataset, which is 10. Considering that the study utilizes CPU as the computing unit, the task_type parameter is set to CPU. Furthermore, since the evaluation metric focuses on the AUC value, the eval_metric parameter is set to AUC.

The ROC curve for scenario 1, 2, 3 are displayed in Figure 7 shows an AUC value of 0,789583, Figure 8 shows an AUC value of 0,80329, and Figure 9 shows an AUC value of 0,781066, respectively, they provided a visual representation of the experimental results.

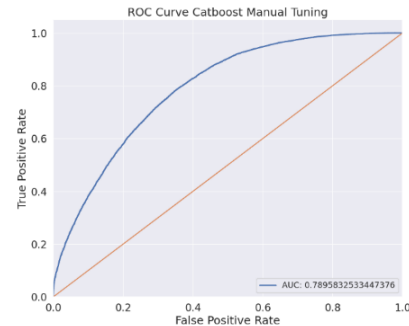


Figure 7. The result of the ROC curve for Scenario 1 (80:20)

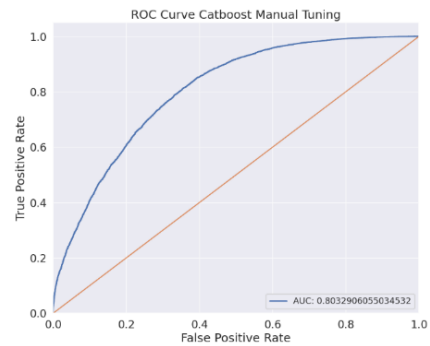


Figure 8. The result of the ROC curve for Scenario 2 (90:10)

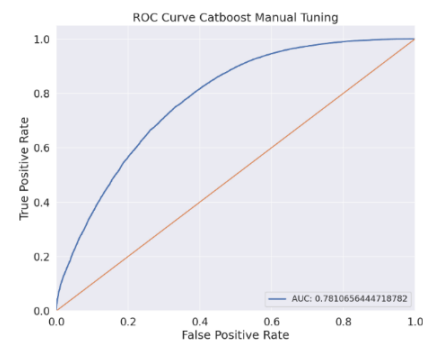


Figure 9. The result of the ROC curve for Scenario 3 (70:30)

Those curves exhibit similar shapes that indicate consistent performance; the blue curve illustrates the performance of a specific classifier or predictive model, plotting the relationship between the true positive rate (sensitivity) and the false positive rate (1-specificity) at various threshold values. In contrast, the orange line represents the center line, which corresponds to the ROC curve of a random classifier with no discriminatory ability. The center line is a straight line connecting the points (0, 0) and (1, 1) on the ROC plot.

Based on the results of Table 5, the scenario 1 experiment yielded an AUC value of 0,789583 with a

training time of 178 seconds. This result means that hyperparameter depth = 8 given around 70% true in classifying credit risk. Meanwhile in scenario 2, though ratio data training:testing 90:10 using the same parameter with scenario 1 obtained an AUC value of 0,803291 through a training time of 255 seconds. Lastly, the scenario 3 experiment found an AUC value of 0,781066 with a training time of 173 seconds. It is caused by the distinguish hyperparameter use, depth = 7 and ratio training:testing 70:30.

Table 5. Result of AUC from 3-Scenario

Method	AUC Score	Training Time (Second)
Scenario 1 (80:20)	0,789583	178
Scenario 2 (90:10)	0,803291	255
Scenario 3 (70:30)	0,781066	173

The results indicate that scenario 2 experiment outperforms in predicting credit risk to support P2P lending compared to the other scenarios in terms of AUC results, with scenario 1 experiment ranked second and the scenario 3 experiment ranked last. However, considering the time in producing model prediction of credit risk, scenario 3 experiment with ratio data training was used smaller than the other exhibits a faster training time compared to the scenario 1 experiment, which is ranked second, while the scenario 2 experiment has the longest training time among the three scenarios.

The results from all three scenarios can be considered favorable as the AUC values obtained in each scenario were close to 1. This proximity to 1 indicates that the model was effective in properly identifying credit risk. Despite achieving satisfactory results, the outcomes obtained by this model still fall short of the best results. This limitation can be attributed to the large number of data records and features, which hinder the optimal performance of the research model. The model's performance could potentially be enhanced by incorporating the feature importance method in the selection of features. Nevertheless, based on the obtained results, it can be concluded that the CatBoost model exhibits effective performance in identifying credit risk. These findings underscore the model's robust capability in detecting credit risk.

4. Conclusion

The aims of this study already succeeded in identifying credit risk in P2P lending according to the experiments result, using 171,257 raw Bondora loan data. It had been proven by three different scenarios that were conducted with data splits of 80:20, 90:10, and 70:30; with the result of AUC values for each scenario were 0.789583, 0.80329, and 0.781066, respectively. Scenario 2 exhibited superior AUC results compared to scenarios 1 and 3. However, the obtained results fell short of optimal outcomes primarily due to the large volume of

data and feature records, as well as the failure to employ the feature importance method. Consequently, the selection of features used in the model remained suboptimal. Moreover, this research concludes that the CatBoost model performs well in detecting credit risk. It is the authors hope that these findings contribute to minimizing credit risk in P2P lending. In future research, it is recommended to improve feature selection compared to the approach used in this study. Alternatively, employing the Principal Component Analysis (PCA) method can be considered.

Acknowledgment

We would like to express our gratitude to Bondora and the Informatics Faculty of Telkom University for their assistance in completing this research project. For this research project can be seen on the link [s17Happiness/TA-Credit-Risk-Detection-In-P2PLending-Using-CatBoost \(github.com\)](https://github.com/s17Happiness/TA-Credit-Risk-Detection-In-P2PLending-Using-CatBoost)

Reference

- [1] P. S. Chong, J. Labadin, and F. Meziane, "Credit Risk Prediction for Peer-To-Peer Lending Platforms: An Explainable Machine Learning Approach," *J. Comput. Soc. Informatics*, vol. 1, no. 2, pp. 1–16, 2022, doi: 10.33736/jcsi.4761.2022.
- [2] D. Li, S. Na, T. Ding, and C. Liu, "Credit risk management of p2p network lending," *Teh. Vjesn.*, vol. 28, no. 4, pp. 1145–1151, 2021, doi: 10.17559/TV-20200210110508.
- [3] L. Zhou, H. Fujita, H. Ding, and R. Ma, "Credit risk modeling on data with two timestamps in peer-to-peer lending by gradient boosting," *Appl. Soft Comput.*, vol. 110, p. 107672, Oct. 2021, doi: 10.1016/J.ASOC.2021.107672.
- [4] M. Bazarbash, "FinTech in Financial Inclusion: Machine Learning Applications in Assessing Credit Risk," *IMF Work. Pap.*, vol. 2019, no. 109, May 2019, doi: 10.5089/9781498314428.001.A001.
- [5] J. Mezei, A. Byanjankar, and M. Heikkilä, "Credit risk evaluation in peer-to-peer lending with linguistic data transformation and supervised learning," *Proc. Annu. Hawaii Int. Conf. Syst. Sci.*, vol. 2018-Janua, pp. 1366–1375, 2018, doi: 10.24251/hicss.2018.169.
- [6] L. Machado and D. Holmer, "Credit risk modelling and prediction: Logistic regression versus machine learning boosting algorithms," 2022.
- [7] N. Nguyen *et al.*, "A Proposed Model for Card Fraud Detection Based on CatBoost and Deep Neural Network," *IEEE Access*, vol. 10, pp. 96852–96861, 2022, doi: 10.1109/ACCESS.2022.3205416.
- [8] X. Li, D. Ergu, D. Zhang, D. Qiu, Y. Cai, and B. Ma, "Prediction of loan default based on multi-model fusion," *Procedia Comput. Sci.*, vol. 199, pp. 757–764, Jan. 2022, doi: 10.1016/J.PROCS.2022.01.094.
- [9] J. D. Turiel and T. Aste, "Peer-to-peer loan acceptance and default prediction with artificial intelligence," *R. Soc. Open Sci.*, vol. 7, no. 6, Jun. 2020, doi: 10.1098/RSOS.191649.
- [10] I. Rahadiyan and N. Mentari, "Keterbukaan Informasi Sebagai Mitigasi Risiko Peer To Peer Lending (Perbandingan Antara Indonesia Dan Amerika Serikat)," *J. Huk. IUS QUIA IUSTUM*, vol. 28, no. 2, pp. 325–347, Jun. 2021, doi: 10.20885/IUSTUM.VOL28.ISS2.ART5.
- [11] R. Hutapea, "Minimalisasi Risiko Kredit (NPL) Pada Fintach Peer to Peer Lending melalui Kewajiban Pelaporan SLIK OJK," *J. Ilm. Mandala Educ.*, vol. 6, no. 2, Oct. 2020, doi: 10.58258/JIME.V6I2.1401.
- [12] W. Li, S. Ding, H. Wang, Y. Chen, and S. Yang,

- "Heterogeneous ensemble learning with feature engineering for default prediction in peer-to-peer lending in China," *World Wide Web*, vol. 23, no. 1, pp. 23–45, Jan. 2020, doi: 10.1007/S11280-019-00676-Y/FIGURES/9.
- [13] Y. Guo, "Credit risk assessment of P2P lending platform towards big data based on BP neural network," *J. Vis. Commun. Image Represent.*, vol. 71, p. 102730, Aug. 2020, doi: 10.1016/J.JVCIR.2019.102730.
- [14] O. Havrylych and M. Verdier, "The Financial Intermediation Role of the P2P Lending Platforms," *Comp. Econ. Stud.*, vol. 60, no. 1, pp. 115–130, Mar. 2018, doi: 10.1057/S41294-017-0045-1/METRICS.
- [15] Y. Song, Y. Wang, X. Ye, D. Wang, Y. Yin, and Y. Wang, "Multi-view ensemble learning based on distance-to-model and adaptive clustering for imbalanced credit risk assessment in P2P lending," *Inf. Sci. (Nij.)*, vol. 525, pp. 182–204, Jul. 2020, doi: 10.1016/J.INS.2020.03.027.
- [16] S. Chen, Q. Wang, and S. Liu, "Credit Risk Prediction in Peer-to-Peer Lending with Ensemble Learning Framework," *Proc. 31st Chinese Control Decis. Conf. CCDC 2019*, no. 1, pp. 4373–4377, 2019, doi: 10.1109/CCDC.2019.8832412.
- [17] Š. Lyócsa, P. Vašaničová, B. Hadji Misheva, and M. D. Vateha, "Default or profit scoring credit systems? Evidence from European and US peer-to-peer lending markets," *Financ. Innov.*, vol. 8, no. 1, pp. 1–21, Dec. 2022, doi: 10.1186/S40854-022-00338-5/FIGURES/3.
- [18] L. Prokhorenkova, G. Gusev, A. Vorobev, A. V. Dorogush, and A. Gulin, "CatBoost: unbiased boosting with categorical features," *Adv. Neural Inf. Process. Syst.*, vol. 31, 2018, Accessed: Nov. 20, 2022. [Online]. Available: <https://github.com/catboost/catboost>
- [19] Y. Zhang, Z. Zhao, and J. Zheng, "CatBoost: A new approach for estimating daily reference crop evapotranspiration in arid and semi-arid regions of Northern China," *J. Hydrol.*, vol. 588, no. May, p. 125087, 2020, doi: 10.1016/j.jhydrol.2020.125087.
- [20] A. V. Dorogush, V. Ershov, and A. Gulin, "CatBoost: gradient boosting with categorical features support," pp. 1–7, 2018, [Online]. Available: <http://arxiv.org/abs/1810.11363>
- [21] J. T. Hancock and T. M. Khoshgoftaar, "CatBoost for big data: an interdisciplinary review," *J. Big Data*, vol. 7, no. 1, 2020, doi: 10.1186/s40537-020-00369-8.
- [22] X. Huang, X. Liu, and Y. Ren, "Enterprise credit risk evaluation based on neural network algorithm," *Cogn. Syst. Res.*, vol. 52, pp. 317–324, Dec. 2018, doi: 10.1016/J.COGSYS.2018.07.023.