

M. Fostyak¹, L. Demkiv²^{1,2}Ivan Franko National University of Lviv, Ukraine
50 Drahomanova St., Lviv, 79005²lidia.demkiv@gmail.com¹<https://orcid.org/0000-0003-4670-5834>²<https://orcid.org/0009-0002-0185-6364>

HYBRID OPTIMIZED BRB–ML MODEL FOR CREDIT RATING PREDICTION IN OPEN BANKING SYSTEMS

Abstract. This paper presents a hybrid decision support system that integrates a Belief Rule Base (BRB) model, a Machine Learning (ML) model, and Particle Swarm Optimization (PSO). The system predicts the credit rating of fintech clients based on data obtained through the Open Banking API. For a real-world dataset, the classification ML model Random Forest demonstrates high predictive accuracy. The BRB model is represented by a complete set of rules covering four attributes with three referential values for each attribute. By analyzing the distribution of the number of activated rules and their activation weights, the representativeness of the data was assessed. A group of rules with low activation weights was identified, corresponding to the referential value *Middle* for one of the attributes. PSO was applied to optimize BRB parameters, ensuring both interpretability of the results and the capability to model incomplete data. For model validation, data were synthesized using a Rule-Based Data Synthetic approach. The results show that BRB models maintain stable accuracy across real and validation datasets. In contrast, the accuracy of the ML model decreased on validation data by an average of 5% compared to the initial dataset. These findings highlight the advantages of the hybrid approach, which combines the accuracy of ML with the interpretability and robustness of BRB in data-limited environments.

Keywords: machine learning, neural networks, fuzzy logic, BRB, ER, open banking, credit rating.

Introduction

In modern research, the use of open data is gaining increasing attention, as it creates new opportunities for building predictive models while simultaneously ensuring the explainability of decision support system outputs. Open data span a wide range of domains, from ecology and transportation to medicine, sociology, education, and finance.

The financial sector has traditionally been characterized by a high degree of opacity and limited access to data. This has complicated the development of flexible and transparent models for risk assessment and decision-making. However, the advancement of digital technologies, the growing role of fintech companies, and the implementation of the open banking concept are significantly transforming approaches to credit scoring, risk forecasting, client reliability assessment, and asset portfolio management. Information on bank accounts and client transactions, obtained with customer consent through the Open Banking API, provides fintech companies – particularly those specializing in small business and consumer lending – with valuable insights.

The transition to an open data economy through open banking is an initiative launched by several governments, including the European Union and the United Kingdom. By October 2021, nearly 80 countries had engaged in such governmental initiatives [1]. On August 1, 2025, Ukraine joined the open banking framework. Studies [1,2] analyze the impact of open banking on the formation of interest rates, default levels, and credit portfolio structures. These works demonstrate that it becomes feasible to extend credit to higher-risk firms with lower credit ratings and elevated default probabilities. Such shifts must be incorporated into credit rating models and credit portfolio optimization strategies.

In [3], machine learning models were developed to predict the probability of inflow or outflow of a client's banking data – i.e., whether a client imports financial data from another institution into the bank or exports it in the opposite direction.

The development of open banking thus exerts a significant influence on the quality of economic indicators in the lending sector. Enhanced access to detailed and standardized data on client income, expenses, and financial behavior creates the necessity of incorporating

these new open data sources into credit rating models to improve their predictive power.

Problem Statement

In the context of modern financial technologies, models capable of combining high accuracy in credit rating prediction with interpretability of results are of particular importance. The aim of this study is to develop an approach that integrates an expert knowledge base (BRB), machine learning (ML) algorithms, and the Particle Swarm Optimization (PSO) algorithm to improve the efficiency of open banking data analysis, particularly in tasks related to credit rating prediction and risk assessment.

Traditional banks possess extensive and long-term information about their clients, covering lengthy periods and providing a more comprehensive view of financial behavior. In contrast, fintech companies typically operate with a smaller customer base and have access only to relatively short-term data, often no more than 90 days. This asymmetry in the depth and duration of data complicates the construction of accurate predictive models, especially in risk assessment. Moreover, the limited number of clients in fintech datasets may fail to capture the full spectrum of behavioral characteristics, leading to biased conclusions. This creates the need for additional research into the completeness of client datasets and decision rules, as well as the development of not only highly accurate ML solutions but also explainable models capable of integrating incomplete data, preserving decision-making transparency, and ensuring trust among both users and regulators.

In [4], the effectiveness of traditional bank credit scoring models was compared with a deep learning model trained solely on transaction data obtained via the Open Banking API. The evaluation was carried out using AUC and Brier metrics. The findings revealed that models relying exclusively on open banking data can be surprisingly effective in default prediction compared to traditional credit scoring approaches. However, the study also reported a decline in the AUC metric when applied to data from new clients.

Such characteristics of data structures give rise to several interrelated challenges that

must be addressed. These include ensuring the representativeness of client samples, integrating incomplete data into decision-making models, and developing approaches that combine high predictive accuracy with transparency and interpretability of results.

To address these issues, this work proposes a hierarchical BRB structure and a hybrid architecture that integrates ML and BRB components. Model parameters are optimized using the Particle Swarm Optimization (PSO) algorithm, which accounts for both the number of activated rules and the weights of individual rules and attributes. This approach ensures system adaptability to different types of input data derived from open banking transactions, enhances predictive accuracy, and simultaneously preserves interpretability in the decision-making process.

Literature Review

In the financial sector, there is a growing interest in the use of open data for developing risk assessment models and decision support systems. An important direction of such research is the optimization of model parameters, particularly for Belief Rule Base (BRB) systems with the Evidential Reasoning (ER) algorithm, as well as hybrid systems that enhance predictive accuracy while preserving result interpretability.

Classical BRB systems are typically constructed by defining referential values of attributes, which in complex tasks may lead to excessively large BRB systems. Hierarchical structures such as HBRB-ER allow the construction of sub-models with smaller rule sets and their subsequent combination. Implementing hierarchical structures reduces the total number of rules and decreases computational load through stepwise aggregation of sub-model results. ER aligns well with hierarchical aggregation since intermediate-level results can be treated as “evidence” for higher levels.

In [5], the accuracy and interpretability of several BRB models, including the hierarchical HBRB model, were compared. The models analyzed the number of activated rules and the distribution of activation weights. For performance evaluation and prediction

comparison, two forecasting models—Random Forest (RF) and Artificial Neural Networks (ANN)—were selected. The comparison was conducted using mean squared error (MSE). Results showed that HBRB improves structural scalability and interpretability of BRB while maintaining model accuracy.

The construction of a hierarchical BRB structure in [6] was implemented to improve performance in imbalanced multiclass classification and to allow operation with any number of classes. The hierarchical BRB contains one primary BRB at the first level and several sub-BRBs at the second level. The study proposes a novel BRB system in a hierarchical structure with XGBoost-based feature selection, simultaneously addressing multiclass imbalance and combinatorial explosion.

In [7], a method for generating and reducing the BRB rule base based on redundancy was proposed. Similar rules are grouped and merged to ensure coverage of all possible situations without generating unnecessary rules.

The study in [8] proposed a BRB construction method based on decision trees and modified the classical BRB so that rule attributes could be expressed as intervals rather than single values. Rule parameters trained with the Differential Evolution (DE) algorithm were optimized and adjusted to further improve system performance. A similar approach of interval-based referential values with subsequent nonlinear optimization for obtaining optimal values was applied in [9].

In [10], both global and local interpretability of BRB were presented. The significance of rules and their attributes in the rule base reflects the global decision-making process of BRB. The importance of rules activated by a given data point, as well as the significance of referential values in the attributes of those rules, provides local interpretability for individual cases.

Rule-based algorithms in hybrid systems play a key role in generating explanations for end-users, thereby ensuring trust in the outcomes of automated analysis. In [11], a fuzzy rule-based classifier was developed

using the metaheuristic Gravitational Search Algorithm (GSA).

To address the combinatorial explosion in BRB expert systems, [12] proposed a new optimization method that considers the number of activated rules. In cases of a large number of activated rules, parallel optimization is applied through decomposition of training data, with activation coefficients used in each sub process.

One promising approach to enhancing decision support accuracy is the partitioning of data or BRB rule sets into separate groups or clusters. This allows for capturing differences between data subsets or rule patterns, balancing datasets while preserving as much information as possible. In [13], Evidential C-Means (ECM—an extension of Fuzzy C-Means) was applied to generate *credal partitions* of data. This method enables objects to belong not only to a single cluster or multiple clusters with fuzzy membership, but also to unions of clusters or uncertainty groups. A systematic method was then developed to construct BRB rules based on the obtained belief distributions. After entropy-based optimization of distributions, a compact BRB was achieved with an improved trade-off between accuracy and interpretability.

The problem of class imbalance can significantly bias system decision-making. To address this, [14] proposed a new Evidential BRB (EBRB) system based on fuzzy C-means clustering (FCM-EBRB). First, FCM clustering was used for oversampling positive cases and under sampling negative ones to achieve balance. Next, the EBRB construction method was refined, and the system was optimized with an efficient parameter learning strategy.

A comprehensive study of hybrid BRB expert systems in [15] demonstrated that the number of publications on this topic has been gradually increasing over the past five years. Such systems are applied in research across various economic domains worldwide.

Data Processing and Transformation

The data obtained through the OpenBank API in JSON format are highly detailed (date, amount, currency, transaction type, and additional metadata) and heterogeneous, as

they include information about transactions, balances, as well as clients' regular incomes and expenses. For subsequent use of these data in hierarchical Belief Rule Base (HBRB) systems and machine learning (ML) algorithms, careful transformation is required.

Client expenditures from the OpenBanking API are aggregated into dozens of categories over a limited time period (e.g., 90 days). Direct use of such a large number of attributes leads to a combinatorial explosion of rules and complicates the interpretability of BRB systems, as the rules become excessively specific. Therefore, thematic aggregation of data was carried out into six groups: periodic and non-periodic income, mandatory, non-mandatory, and risky expenditures, as well as the difference between income and expenses. To construct a BRB model with a complete set of rules, further reduction was applied: instead of six initial groups, four integrated categories were formed. Relative indicators (shares) were used to reflect the contribution of each group within the overall income structure.

As a result, the following categories were defined as input parameters for the BRB:

- Stability of Income (SI): determines the regularity of account inflows and is defined as the ratio of periodic income to total income.
- Stability of Discretionary Spending (SDS): indicates the share of income that remains stable after covering non-mandatory expenses.
- Stability of Risky Spending (SRS): indicates the share of income that remains stable after covering risky expenditures.
- Relative Difference (DiFF): defines the proportion of income remaining after all expenses are made.

Correlation matrix analysis shows weak or very weak correlation coefficients for both integrated and non-integrated datasets [16–18]. Aggregating the initial data for 273 clients into fewer groups results in the loss of detailed information regarding client behavioral characteristics. To account for these behavioral patterns, Figure 1 presents the structure of the hierarchical BRB model. At the upper level, integrated categories are analyzed, while at the lower level, more detailed subgroups of data related to non-mandatory and risky expenditures are assessed. These categories are

particularly important for evaluating financial stability: unlike mandatory expenditures, non-mandatory spending may be reallocated for loan repayment, while risky expenditures must be examined in detail as a key parameter of unstable financial behavior.

In the hierarchical BRB model, the results of rule processing at the lower level are aggregated and represented as generalized indicators at the higher level. Thus, the information is incorporated in the form of integrated belief values for the higher-level categories.

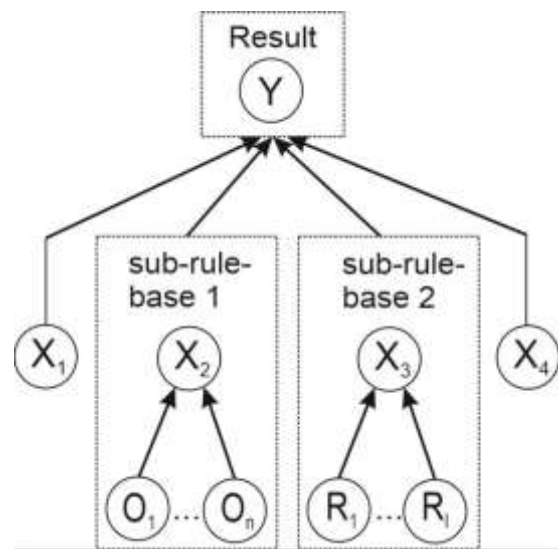


Fig. 1. Structure of the hierarchical BRB model

When working with open banking data, it is necessary to account for the limited number of clients in the sample. This does not always guarantee full representativeness of behavioral characteristics. To overcome this limitation, the Belief Rule Base (BRB) methodology with a complete rule base is applied. BRB enables handling incomplete and uncertain information and allows interpolation of conclusions even for underrepresented client behavior scenarios.

Rule Base Construction in Belief Rule-Based Systems

This section discusses the application of the Belief Rule Base (BRB) model in combination with the Evidential Reasoning (ER) algorithm for determining clients' credit ratings. The BRB model formalizes knowledge in the form of a set of rules with belief distributions, while the ER algorithm is used to

aggregate these rules and compute the final rating outcome. Each rule in BRB is represented as follows:

$$R_k: IF(x_1 \text{ is } A_{1n}) \wedge (x_2 \text{ is } A_{2n}) \wedge \dots \wedge (x_i \text{ is } A_{in})$$

then (y is F_j) with belief degrees β_{kj}

with rule weight Θ_k

with attribute weight $\delta_1 \dots \delta_n$

where x_i are the input parameters and A_{in} are their reference values. Each input parameter (attribute) is discretized into a limited number of reference levels. The reference values, which represent the main quantitative states of each attribute, are presented in Table 1.

Table 1. Reference values for the initial data

	$A_{i,1}$ (Low)	$A_{i,2}$ (Middle)	$A_{i,3}$ (High)
Stability of Income (SI)	30%	60%	90%
Share of Discretionary Expenditures (SDS)	45%	65%	80%
Share of Risky Expenditures (SRS)	80%	90%	95%
Financial Savings (DiFF)	10%	18%	21%

The resulting value of the output variable, namely the credit rating, is determined by the degree of belief β_{kj} assigned to each of the two categories of the output variable: F_1 = Fail, F_2 =Success.

To implement the calculation of belief degrees using the BRB+ER method, the following stages of data processing must be carried out: calculation of the matching degree, computation of the activated rule weights (Activation Weight), aggregation of beliefs using the Evidential Reasoning (ER) algorithm, and computation of the final belief distribution for the credit rating or the expected utility for further interpretation and comparison of the credit rating with machine learning results. Since all input parameters have quantitative representations, the matching degree is calculated by linear interpolation between the corresponding reference values.

The activation weight of a rule is defined as the product of all matching degrees for the attributes included in the rule and the rule weight. Equation (1) for normalized activation weights incorporates the rule weight Θ_k and the normalized weight of the rule's attribute $\underline{\delta}_i$, which is computed according to Equation (2).

$$\omega_k = \frac{\Theta_k \prod_{i=1}^I (\alpha_{i,n}^k)^{\delta_i}}{\sum_{k=1}^K \Theta_k \prod_{i=1}^I (\alpha_{i,n}^k)^{\delta_i}} \quad (1)$$

$$\underline{\delta}_i = \frac{\delta_i}{\max_{j=1,2,\dots,N} \{\delta_j\}} \quad (2)$$

In BRB+ER systems, parameter aggregation is performed simultaneously for all rules using specialized aggregation formulas that take into account normalized rule weights and their belief degrees. This allows information from all rules to be combined consistently, properly considering both the strength of rule activation and potential conflicts among them. The final result is expressed by Equation (3) as an integrated belief vector, which reflects the overall evaluation of all rules with respect to each possible output state. The conflict coefficient μ ensures proper scaling of the final belief degrees, maintaining balance between the information provided by the rules and the system's level of uncertainty. The value of μ decreases in cases of significant conflict among rules or high uncertainty, while in cases of consistent and well-defined data, it approaches unity.

$$\beta_m = \frac{\mu \left[\prod_{k=1}^K (\omega_k \beta_{n,k} + 1 - \lambda) - \prod_{k=1}^K (1 - \lambda) \right]}{1 - \mu \left[\prod_{k=1}^K (1 - \omega_k) \right]} \quad (3)$$

$$\mu = \left[\sum_{n=1}^N \prod_{k=1}^K (\omega_k \beta_{n,k} + 1 - \lambda) - (N-1) \prod_{k=1}^K (1 - \lambda) \right]^{-1}$$

$$\lambda = \omega_k \sum_{n=1}^N \beta_{n,k}$$

$$Predict = \sum_{i=1}^N u_i \beta_i$$

where u_i (utility) is the numerical representation of the reference level of the output attribute.

Table 2 presents the initial configuration of the rule base prior to applying the PSO optimization algorithm. For each rule, the following are specified: the rule number, the group to which the rule belongs, the combination of input attributes and their levels. Table 2 also provides the initial parameter values before optimization. The rule weights are set to unity, and each of the four attributes initially has an evenly distributed weight of 0.25. The belief degrees distribute values between the possible outcomes F_1 = Fail and F_2 =Success.

All values shown in Table 2 are used as the starting configuration for subsequent optimization, where the rule weights and belief degrees are adjusted using the PSO algorithm to minimize error.

Hybrid Integration of BRB, Machine Learning and PSO Optimization

Figure 2 shows the architecture of the hybrid model. The hybrid model consists of four interconnected components: BRB, ML, PSO, and MSE.

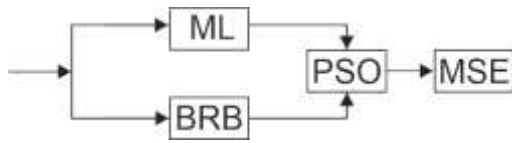


Fig. 2. Hybrid model architecture

Categorized data obtained from the Open Banking API are fed into the system as inputs. The data are simultaneously directed to the machine learning (ML) block and the belief rule base (BRB) block with belief degrees. Optimization in the hybrid system is performed to adjust the weights of rules and attributes in the BRB so that its outputs align with the predictions of the ML model while minimizing the mean squared error (MSE). This improves accuracy while preserving system transparency.

Such an architecture combines the strong predictive capability of ML with the transparency and interpretability of BRB, while also enabling a deeper analysis of the distribution of activated rules to assess the representativeness of client data.

$$J(p) = \frac{1}{N} \sum_{j=1}^N (y_j^{BRB} - y_j^{true})^2. \quad (6)$$

Before optimization, so-called “dead rules” were removed from the rule base. These rules describe attribute combinations that do not occur in real data (for example, “very low income,” “very high discretionary expenses,” and “high difference between income and expenses” in the financial behavior model). For such rules, the matching degree equals zero, meaning they are never included in the ER aggregation process. Eliminating dead rules reduces the dimensionality of the task and accelerates optimization.

In the Particle Swarm Optimization (PSO) algorithm, each possible solution to the problem is modeled as a particle in a multidimensional parameter space. Each particle is characterized by its position and velocity, which determine its movement within the search space. The movement of particles is influenced by a combination of the particle’s own experience (pbest) and the collective experience of the entire swarm (gbest). Through this gradual exchange of information and adjustment of positions, the swarm converges toward an optimal or near-optimal solution. The velocity update is calculated according to the following formula:

$$v_i(t+1) = \omega v_i(t) + c_1 r_1 (pbest_i - x_i(t)) + c_2 r_2 (gbest_i - x_i(t)) \quad (4)$$

The inertia coefficient ω was linearly decreased from 0.9 to 0.4 throughout the optimization process. The parameters were set as $c_1=c_2=1.6$, $r_1=r_2=0.8$, and the particle velocity was limited to 20% of the parameter range. The position update is given by:

$$x_i(t+1) = x_i(t) + v_i(t+1) \quad (5)$$

where x_i is the position of a particle, corresponding to the parameter values p : belief degrees, rule weights, and attribute weights for the attributes SI, SDS, SRS, and DiFF. To ensure model correctness, the belief degrees were normalized to sum to 1, and the attribute weights were scaled so that their total value equaled 1. The stopping criterion was defined as the stabilization of the objective function:

Table 2. Non-optimized parameters of the rule base

Rule Group	Rule Number	Rule Weight	Rule				Attribute weight				Belief Degree (F_1, F_2)
			SI	SDS	SRS	DiFF	SI	SDS	SRS	DiFF	
1 (1-27)	19	1	L	H	L	L	0.25	0.25	0.25	0.25	0.875, 0.125
	...	...	...	...	...	...	...	...	...	...	...
2 (28-54)	28	1	M	L	L	L	0.25	0.25	0.25	0.25	0.750, 0.250
	...	...	...	...	...	...	...	...	...	...	...
3 (55-81)	68	1	H	M	M	M	0.25	0.25	0.25	0.25	0.275, 0.725
	...	...	...	...	...	...	...	...	...	...	...

Table 3. Optimized parameters of the rule base

Rule Group	Rule Number	Rule Weight	Rule				Attribute weight				Belief Degree(F_1, F_2)
			SI	SDS	SRS	DiFF	SI	SDS	SRS	DiFF	
1 (1-27)	19	0.85	L	H	L	L	0.23	0.30	0.19	0.28	0.952, 0.048
	...	...	...	...	...	...	...	...	...	...	...
2 (28-54)	28	0.69	M	L	L	L	0.24	0.26	0.22	0.29	0.731, 0.269
	...	...	...	...	...	...	...	...	...	...	...
3 (55-81)	68	0,97	H	M	L	M	0.16	0.26	0.27	0.31	0.154, 0.846
	...	...	...	...	...	...	...	...	...	...	...

Figure 3 shows the dependence of the mean squared error (MSE) on the number of PSO iterations. The high dimensionality of the parameter space $\mathbf{p}$ in the optimization task results in a slow change of MSE during the first hundred iterations. The particles are distributed relatively widely, and the algorithm requires significant time to approach the region of correct solutions.

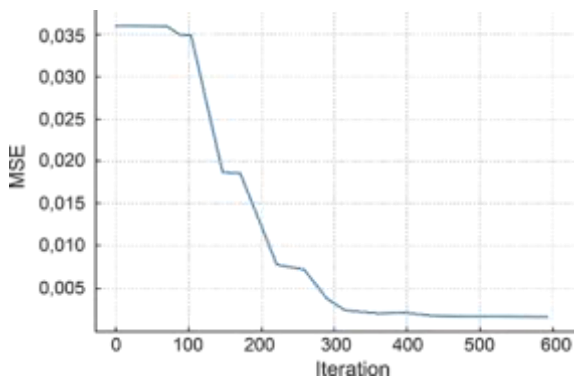


Fig. 3. Convergence curve of the optimization process: dependence of Mean Squared Error (MSE) on the number of iterations

The two small plateaus observed on the MSE convergence curve indicate temporary stabilization of the PSO algorithm in intermediate local minima, possibly related to the presence of several client groups in the data. Due to the stochastic components and inertia, some particles escape from local minima, after which the error continues to

decrease toward the global minimum. The numerical results of the model parameters after optimization are shown in Table 3.

The attribute weights are optimized separately for each rule; therefore, the same attribute may be highly significant in one scenario but have almost no influence in another. A local weight plot for a single rule illustrates how attributes “compete” for importance within that specific rule, while overall balance is maintained through normalization. Figure 4 shows the distribution of attribute weights after optimization for Rule 68. The graph demonstrates that, for Rule 68, income stability proved to be a less significant factor for the predicted outcome compared to the influence of discretionary spending (SDS), risky spending (SRS), and their relative difference (DiFF), contrary to the initial assumption.

The analysis of rule activation weights showed that for the attribute SI with the referential value *Middle*, a portion of the rules in the BRB system were activated with low weights. This indicates that the available client data lack sufficient examples where SI takes on an intermediate level. Consequently, rules corresponding to combinations with $SI = Middle$ remain latent in the system, preserving readiness for generalization in the event that new clients with such characteristics appear. The low activation of these rules also points to

asymmetry in the data distribution: in the current dataset, clients with extreme SI values (*Low* or *High*) dominate, while intermediate profiles occur much less frequently.

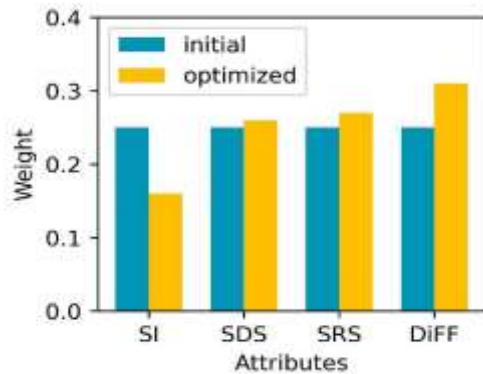


Fig. 4. Attribute weights adjustment for Rule 68 before and after optimization

To evaluate the effectiveness of the proposed hybrid system, a synthetic dataset was generated using the Rule-Based Data Synthetic approach. This dataset was constructed to reflect all client behavioral patterns encoded in the BRB rule base. Table 4 presents a comparison of the accuracy of Random Forest and the optimized BRB+ER system on both initial and synthetic data.

Table 4. Accuracy of ML and BRB-ER models on real and synthetic data

	Random Forest	BRB+ER optimized
Initial data	0.93±0.02	0.92±0.01
Synthetic data	0.88±0.02	0.91±0.01

On the initial data, both models demonstrate comparable levels of accuracy, which indicates that the optimized BRB+ER system achieves nearly the same performance as a classical machine learning algorithm. The decline in Random Forest performance can be explained by the fact that incomplete initial data limited its ability to generalize to new scenarios. In contrast, the BRB+ER system, owing to its rule structure and evidential aggregation mechanism, exhibits greater stability and robustness to changes in the input data.

Thus, it has been shown that the optimized BRB-ER system maintains a high level of accuracy and interpretability when applied to synthetic data. This makes it suitable

for practical applications where available data are incomplete or contain uncertainty.

Conclusions

The developed system integrates BRB rules with ML blocks to improve predictive accuracy while preserving interpretability and applies the Particle Swarm Optimization (PSO) algorithm for fine-tuning belief degrees, rule weights, and attribute weights. During BRB operation, rules with low activation weights are triggered due to the absence of clients in the dataset with the corresponding attribute combinations. This indicates limited data coverage, meaning that certain classes or attribute combinations are underrepresented in the original sample.

When transitioning to synthetically generated data, which reflect the full space of possible input scenarios, the accuracy of the ML model decreases due to insufficient training on similar examples, whereas the accuracy of the optimized BRB-ER model remains virtually unchanged. This behavior emphasizes the importance of maintaining a complete rule base to enable generalization to new scenarios, even if the current dataset does not include examples for all conditions.

Thus, the proposed approach not only supports the construction of more robust decision-making models but also enables the analysis of activated rule counts and the contribution of individual attributes—an aspect particularly critical in the financial sector, where transparency and explainability are of paramount importance.

References

1. Zhiguo He, Jing Huang, Jidong Zhou (2023). Open banking: Credit market competition when borrowers own the data, *Journal of Financial Economics*, Volume 147, Issue 2, <https://doi.org/10.1016/j.jfineco.2022.12.003>
2. Xiong Rui (2024) Open banking and fintech lending: evidence from a crowdfunding platform. <http://dx.doi.org/10.2139/ssrn.5046894>
3. João B. G. de Brito, Rodrigo Heldt (2025). Predicting and Explaining Customer Data Sharing in the Open Banking arXiv:2507.01987 <https://doi.org/10.48550/arXiv.2507.01987>
4. Hjelkrem, L. O., de Lange, P. E., & Nettet, E. (2022). The Value of Open Banking Data for Application Credit Scoring: Case Study of a Norwegian

Bank. Journal of Risk and Financial Management, 15(12), 597.

<https://doi.org/10.3390/jrfm15120597>

5. Xiuxian Yin, Xin Zhang, Hongyu Li, Yujia Chen, Wei He (2023) An interpretable model for stock price movement prediction based on the hierarchical belief rule base, *Heliyon*, Volume 9, Issue 6, <https://doi.org/10.1016/j.heliyon.2023.e16589>

6. Guanxiang Hu, Wei He, Chao Sun, Hailong Zhu, Kangle Li, Li Jiang (2023) Hierarchical belief rule-based model for imbalanced multi-classification, *Expert Systems with Applications*, Volume 216, <https://doi.org/10.1016/j.eswa.2022.119451>

7. Fei Gao, Wenhao Bi, (2023) A fast belief rule base generation and reduction method for classification problems, *International Journal of Approximate Reasoning*, Volume 160, <https://doi.org/10.1016/j.ijar.2023.108964>.

8. Y. Fu, Z. Yin, M. Su, Y. Wu, G. Liu, (2020) Construction and Reasoning Approach of Belief Rule-Base for Classification Base on Decision Tree in *IEEE Access*, vol. 8, pp. 138046-138057, doi: 10.1109/ACCESS.2020.3012453

9. Sun, C., Yang, R., He, W. et al. A novel belief rule base expert system with interval-valued references. *Sci Rep* 12, 6786 (2022). <https://doi.org/10.1038/s41598-022-10636-8>

10. Sachan, S., Yang, J.-B., & Xu, D.-L. (2020). Global and Local Interpretability of Belief Rule Base. In *Proceedings of FLINS2020* World Scientific Publishing Co. https://doi.org/10.1142/9789811223334_0009

11. R. Kumar et al., (2024) Practicing the Rule-Based Classifier Using Metaheuristics Gravitational Search Algorithm to Generate a Credit Score Model to Avail Loans, 1st International Conference on Sustainable Computing and Integrated Communication in Changing Landscape of AI, Greater Noida, India, pp. 1-9, doi: 10.1109/ICSCAI61790.2024.10866544

12. Xiang, G., Wang, J., Han, X. et al. A novel optimization method for belief rule base expert system with activation rate. *Sci Rep* 13, 584 (2023). <https://doi.org/10.1038/s41598-023-27498-3>

13. Jiao, L., Geng, X., & Pan, Q. (2019). Compact Belief Rule Base Learning for Classification with Evidential Clustering. *Entropy*, 21(5), 443. <https://doi.org/10.3390/e21050443>

14. Yang-Geng Fu, Ji-Feng Ye, Ze-Feng Yin, Long-Jiang Chen, Ying-Ming Wang, Geng-Geng Liu, (2021) Construction of EBRB classifier for imbalanced data based on Fuzzy C-Means clustering, *Knowledge-Based Systems*, Volume 234, 107590, <https://doi.org/10.1016/j.knosys.2021.107590>.

15. K. Derrick (2024) A Comprehensive Survey of Belief Rule Base (BRB) Hybrid Expert system: Bridging Decision Science and Professional Services *arXiv:2402.16651* <https://doi.org/10.48550/arXiv.2402.16651>

16. Fostyak M., Demkiv L. (2024) Development of data mesh data platform with ml domain of data analysis, *Electronics and information technologies*. 2024. 27. doi: <https://doi.org/10.30970/eli>

17. M. Fostyak, L. Demkiv (2025) A data-centric approach to building ai models for determining the credit rating of fintech company clients based on open banking // *ISSN 2710 – 1673 Artificial Intelligence 2025 № 1*. <https://doi.org/10.15407/jai2025.01.132>

18. Fostyak M., (2024) Development of an AI domain in a data mesh network for customer credit classification using transaction data, *IEEE 19th International Conference on Computer Science and Information Technologies*. DOI:10.1109/CSIT65290.2024.10982569

19. T. M. Shami, A. A. El-Saleh, M. Alswaitti, Q. Al-Tashi, M. A. Summakieh and S. Mirjalili (2022) Particle Swarm Optimization: A Comprehensive Survey, in *IEEE Access*, vol. 10, pp. 10031-10061, doi: 10.1109/ACCESS.2022.3142859. r

The article has been sent to the editors 25.08.25.

After processing 11.09.25.

Submitted for printing 30.09.25.

Copyright under license CCBY-SA4.0.