

INTERPRETABLE MACHINE LEARNING APPROACHES FOR STUDENT DROPOUT PREDICTION IN HIGHER EDUCATION

Arūnas MINCEVIČIUS[✉]

Department of Information Technologies, Faculty of Fundamental Sciences, Vilnius Gediminas Technical University, Vilnius, Lithuania

Article History:

- received 25 May 2026
- accepted 15 June 2026

Abstract. Student dropout remains one of the most significant challenges in higher education, affecting academic performance, financial sustainability, and strategic planning within universities. This study presents an approach to predicting student dropout risk using machine learning methods and educational analytics. The research is based on an open-access dataset from the UCI Machine Learning Repository containing 4,424 records and 36 attributes without missing values. Three experimental datasets were constructed using different preprocessing strategies, including class balancing, logarithmic feature transformation, and processing of the Enrolled category. Logistic Regression (LR), Support Vector Machine (SVM), Decision Tree (DT), Random Forest (RF), k-Nearest Neighbours (k-NN), and Neural Network (NN) models were implemented and compared. The results showed that Logistic Regression and Random Forest achieved the highest performance, with accuracy above 90% and ROC AUC values of up to 0.96. The most significant risk factors included academic performance during the first and second semesters, tuition fee payment status, outstanding debt, scholarship status, age at enrolment, and study programme. The findings indicate that machine learning methods can effectively support early dropout prediction and decision-support systems in higher education institutions.

Keywords: educational analytics, educational data mining, machine learning, student dropout prediction, classification, logistic regression, random forest, SHAP.

[✉]Corresponding author. E-mail: arunas.mincevicius@gmail.com

1. Introduction

Student dropout remains one of the most significant challenges in higher education. Student dropout has been studied for several decades and remains an important topic in educational research. Early theoretical models emphasized the role of academic and social integration in student persistence, while more recent studies have highlighted the influence of academic performance, financial conditions, motivation, and institutional support on student retention (Spady, 1971; Tinto, 1975; Behr et al., 2021). Withdrawal or dismissal from university may result from a wide range of factors, including academic difficulties, financial constraints, low motivation, insufficient social integration, family-related issues, inappropriate program selection, or the lack of timely institutional support. For students, discontinuing their studies often leads to the loss of time, financial resources, and professional development opportunities. For universities, high dropout rates reduce the efficiency of the educational process, negatively affect institutional performance indicators, decrease financial sustainability, and require the reconsideration of student support strategies.

Traditional approaches to reducing student dropout are typically based on academic advising, mentoring, administrative monitoring, and student counseling. However, such measures are often implemented only after students have already accumulated significant academic or personal difficulties. Consequently, the early prediction of dropout risk has become an increasingly important research problem. If universities are able to identify at-risk students in advance, they can provide timely academic, financial, or psychological support.

The rapid development of information systems in education has created new opportunities for addressing this challenge. Universities accumulate large volumes of student-related data, including admission records, grades, attendance information, completed credits, financial status, study programs, demographic characteristics, and other indicators. These data can be used not only for administrative purposes but also for the development of predictive models. In this context, Educational Data Mining (EDM) plays a particularly important role, as it focuses on discovering patterns in educational data and supporting data-driven decision-making.

The research is based on an open-access dataset obtained from the UCI Machine Learning Repository (<https://archive.ics.uci.edu/>), containing 4,424 records and 36 attributes without missing values. The dataset includes demographic, academic, financial, socio-economic, and macroeconomic characteristics of students. The target variable represents student academic status and consists of three categories: Dropout, Enrolled, and Graduate.

Several experimental datasets were constructed for the study. In the first scenario, the Enrolled category was removed, allowing the problem to be formulated as a binary classification task between Dropout and Graduate. In the second experimental scenario, Enrolled instances were predicted using a neural network model and transformed into one of the two final classes. Additionally, a dataset containing logarithmically transformed numerical features was examined in order to evaluate the impact of different preprocessing approaches on classification performance.

The objective of this study is to evaluate the applicability of interpretable machine learning methods for student dropout prediction and to identify the factors that contribute most strongly to dropout risk. The study compares different data preprocessing strategies and classification models in order to determine which approaches provide the most reliable prediction performance.

The study provides a comparative analysis of machine learning methods for predicting student dropout risk. The research includes exploratory data analysis, feature significance evaluation, and the construction of multiple experimental datasets using different preprocessing strategies. Particular attention is given to the impact of class balancing, logarithmic feature transformation, and feature selection on classification quality. Within the study, the following models were implemented and compared: Logistic Regression (LR), Support Vector Machine (SVM), Decision Tree (DT), Random Forest (RF), k-Nearest Neighbors (k-NN), and Neural Network (NN). In addition, the results are interpreted using statistical analysis and model interpretation techniques in order to identify the factors that most significantly influence the probability of student dropout.

The proposed approach may support the development of Early Warning Systems in higher education institutions, enabling the timely identification of at-risk students and supporting decision-making processes related to student retention.

The remainder of this paper is organized as follows. Section 2 presents related studies on student dropout prediction. Section 3 describes the proposed framework, data preprocessing, machine learning methods, and evaluation strategy. Section 4 presents the experimental results and discussion. Section 5 provides the conclusions and study limitations.

2. Related works

Student dropout research has a long history and encompasses both classical sociological approaches and modern methods of educational analytics and machine learning. Among the most influential theoretical foundations of student attrition research are the works of Tinto (1975, 1987), who proposed the theory of academic and social integration. According to this concept, the probability of successful degree completion is closely associated with the extent of a student's integration into the academic and social environment of the university. Similar ideas were presented by Spady (1971), who examined the role of social adaptation in students' academic persistence. Significant contributions to the development of student attrition theories were also made by Bean (1980), Astin (1993), Terenzini et al. (1994), and Pascarella and Terenzini (2005), who emphasized the importance of academic involvement, interaction with faculty, and satisfaction with the educational environment for student retention.

The development of Educational Data Mining (EDM) and Learning Analytics has substantially expanded traditional approaches to student dropout research. Educational Data Mining is regarded as a field focused on the application of intelligent data analysis techniques to educational processes in order to identify learning patterns, predict academic outcomes, and support decision-making (Baker & Yacef, 2009). Romero and Ventura (2010) note that EDM methods enable the automatic analysis of large volumes of educational data and facilitate the identification of patterns related to student behavior and learning processes. Suhrman et al. (2014) emphasize that educational data mining includes stages of data collection, preprocessing, modeling, and interpretation of results; however, unlike traditional data mining, it focuses specifically on educational data, including academic performance, attendance, activity in learning management systems, and patterns of interaction within the educational environment.

Contemporary studies demonstrate that student dropout is a multifactorial phenomenon influenced by academic, social, demographic, and economic factors. Kehm et al. (2020) identify key groups of factors associated with academic integration, prior academic achievement, personal characteristics, and socioeconomic background. Behr et al. (2021) extend this classification by including financial difficulties, low motivation, family circumstances, and dissatisfaction with the organization of the educational process. Collectively, these studies indicate the complex nature of student attrition and the need for multidimensional analytical approaches. Additional research highlights the importance of financial support and employment status. In particular, Cabrera et al. (1992) demonstrated that financial support contributes to reducing

dropout risk, whereas R. Stinebrickner and T. Stinebrickner (2008) found that balancing work and study negatively affects academic performance and student retention.

One of the most consistent directions in the literature concerns the analysis of academic indicators as predictors of dropout risk. Larsen et al. (2013) demonstrated that secondary school performance and entrance examination results are associated with the likelihood of successful university completion. Similar findings were reported by Aulck et al. (2017), who analyzed an institutional dataset from the University of Washington containing information on approximately 32,500 students. The authors identified the year of admission, birth year, and grades in mathematics, English, chemistry, and psychology as important early indicators of dropout risk. Miguéis et al. (2018), using Decision Tree, Naive Bayes, Support Vector Machine, and Random Forest methods, demonstrated the high predictive significance of entrance examination scores and previous academic performance. The best results were achieved using Random Forest, where accuracy reached 96.1%, indicating the strong potential of ensemble methods for analyzing structured educational datasets. In the study by Orozco and Niguidula (2018), based on academic and demographic data from first-year students, model accuracy was approximately 80%, while GPA was identified as one of the most informative predictors, whereas gender demonstrated a less pronounced effect.

A substantial proportion of recent studies focuses on the application of machine learning methods for the early identification of at-risk students. The most commonly used approaches include logistic regression, Random Forest, Support Vector Machine, Decision Tree, and ensemble methods. Del Bonifro et al. (2020), analyzing data from more than 15,000 students, demonstrated considerable variability in predictive performance depending on educational programs and dataset structure, with model accuracy ranging from 44% to 99%. Such variation suggests a strong dependence of model quality on institutional context, feature composition, and dataset characteristics, thereby limiting the transferability of predictive models across universities. The growing interest in very-early prediction is also reflected in the work of Carballo-Mendivil et al. (2025), who developed an XGBoost-based early warning system using exclusively pre-enrollment data from nearly 40,000 first-year students. Their findings indicate that meaningful dropout risk estimation can be achieved even before the beginning of academic studies.

Naseem et al. (2019) demonstrated the effectiveness of ensemble approaches and the Boruta algorithm for feature selection. Their models achieved an accuracy of approximately 81%, while GPA, gender, year of admission, and educational program were identified as significant predictors. Jiang et al. (2014) demonstrated that early learner behaviour, assessment performance, and social interaction data can be effectively used for academic performance prediction. Using logistic regression models and 10-fold cross-validation, the authors achieved prediction accuracy of up to 92.6%.

In recent years, increasing attention has been directed toward cognitive, psychological, and behavioral factors. Gallego et al. (2021) investigated the influence of motivation, metacognitive strategies, and emotional characteristics, identifying groups of students with different academic risk profiles using cluster analysis. Their findings suggest that low motivation and underdeveloped metacognitive skills are associated with a higher probability of academic difficulties.

Another important research direction concerns online education, learning analytics, and Massive Open Online Courses (MOOCs). Jiang et al. (2014) analysed a four-week biology MOOC hosted on the Coursera platform and demonstrated that first-week assessment performance and participation in discussion forums can be used to predict course completion and certificate attainment. Their findings highlight the importance of early behavioural indicators and learner engagement in online learning environments. More recent learning analytics studies increasingly rely on Learning Management System (LMS) interaction data for early-risk detection. Marcolino et al. (2025), using Moodle log data together with CatBoost optimization techniques, demonstrated that behavioural activity patterns can effectively support the early identification of at-risk students, particularly when combined with imbalance-handling strategies and hyperparameter optimization.

Studies by Khan and Al Zubaidy (2017) demonstrated the potential of multiple linear regression for predicting student academic performance using aptitude test scores, physical training indicators, and learning activity measures. Their findings suggest that early identification of academically at-risk students is possible even during the initial stages of study, supporting the development of timely intervention strategies.

Alongside traditional supervised learning methods, modern studies increasingly employ neural networks and Explainable AI approaches. Zhao et al. (2018) note that convolutional and recurrent neural networks are capable of modeling complex nonlinear relationships in educational data; however, Goodfellow et al. (2016) emphasize the strong dependence of such models on large training datasets and careful hyperparameter tuning. Simultaneously, growing attention is being paid to the interpretability of machine learning models. One of the most widely used Explainable AI methods is SHAP (SHapley Additive exPlanations), which enables the evaluation of individual feature contributions to model predictions. Lundberg and Lee (2017) demonstrated that SHAP provides a unified framework for interpreting complex machine learning models and can be effectively applied to the analysis of academic risk factors. Contemporary studies also emphasize ethical issues related to educational analytics. Slade and Prinsloo (2013) highlight the importance of algorithm transparency, protection of students' personal data, and prevention of discrimination in the use of predictive analytics in education.

Overall, contemporary research demonstrates a growing interest in student dropout prediction at the intersection of educational analytics, Educational Data Mining, and machine learning. Existing studies highlight the significant role of academic, financial, and behavioral factors in assessing dropout risk, as well as the widespread application of supervised learning methods for the early identification of at-risk students. At the same time, research findings strongly depend on dataset characteristics, institutional context, and feature structure, which limits the transferability and generalizability of predictive models across universities. In recent years, particular attention has been devoted not only to improving predictive performance but also to enhancing model interpretability, algorithm transparency, and ethical considerations in educational analytics. These trends indicate the need for the development of more robust, interpretable, and adaptable approaches to student dropout prediction. Table 1 summarizes selected previous studies on student dropout prediction.

Table 1. Summary of selected studies on student dropout prediction

Study	Main contribution	Method(s)	Limitation
Aulck et al. (2017)	Student dropout prediction using academic and demographic data	LR, RF	Moderate predictive performance
Miguéis et al. (2018)	Early identification of at-risk students using predictive modelling	NB, DT, RF	Institution-specific dataset
Naseem et al. (2019)	Feature selection and dropout prediction in computing education	Ensemble DT, Boruta	Evaluated within a single institutional context
Jiang et al. (2014)	Early prediction using learner behaviour and social interaction data	LR	MOOC-specific environment
Marcolino et al. (2025)	LMS-based early-risk detection using Moodle logs	CatBoost	Dependent on LMS activity data
Carballo-Mendivil et al. (2025)	Pre-enrollment dropout prediction and early warning systems	XGBoost	Uses only pre-enrollment information

3. Student dropout prediction framework

3.1. Research framework and data preprocessing

This study employed an educational data mining framework for student dropout prediction using machine learning methods and interpretable artificial intelligence approaches. The research workflow included data preprocessing, feature engineering, class balancing, model training, comparative evaluation, and interpretation of prediction results.

The study utilised an open-access dataset obtained from the UCI Machine Learning Repository. The dataset consisted of 4,424 student records and 36 attributes without missing values. The variables included demographic, academic, financial, socio-economic, and macroeconomic indicators related to student performance and educational status.

Since the target variable contained three categories (Dropout, Enrolled, and Graduate), several experimental dataset configurations were constructed in order to evaluate different preprocessing strategies and classification scenarios. The first dataset configuration focused on binary classification between Dropout and Graduate students. An additional experimental configuration included logarithmic transformation of numerical variables, while another scenario retained the original dataset structure and additionally processed the Enrolled category using neural network-based classification.

Exploratory data analysis and statistical preprocessing were performed prior to model training. Because several numerical variables exhibited non-normal distributions, logarithmic transformation was applied to selected features before model development. Min-max normalization was additionally used to scale numerical variables into a common range, particularly for models sensitive to feature magnitude.

To reduce feature redundancy and improve computational efficiency, correlation analysis was applied during feature engineering. Highly correlated variables were filtered from the training dataset before model development. Additionally, feature importance analysis was later performed using the SHAP (SHapley Additive Explanations) method in order to improve model interpretability.

As the dataset exhibited class imbalance, balancing techniques were applied during the experimental phase. Synthetic Minority Over-sampling Technique (SMOTE) was applied in selected experiments in order to improve class distribution during model training. Figure 1 presents the proposed student dropout prediction framework used in this study.

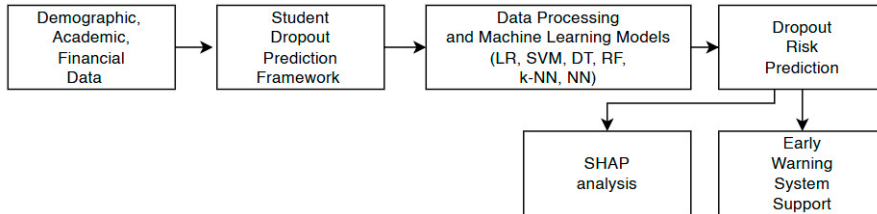


Figure 1. Proposed student dropout prediction framework

The framework combines demographic, academic, and financial data with data processing and machine learning models to predict student dropout risk. The prediction results are further interpreted using SHAP analysis and may support the development of Early Warning Systems in higher education institutions.

3.2. Machine learning framework

Several supervised machine learning algorithms were implemented and compared for student dropout prediction. The experimental framework included Logistic Regression (LR), Support Vector Machine (SVM), Decision Tree (DT), Random Forest (RF), k-Nearest Neighbours (k-NN), and Neural Network (NN) models.

Logistic Regression was used as an interpretable baseline classifier, while Random Forest and Support Vector Machine models were employed to evaluate nonlinear classification performance. Decision Tree and k-Nearest Neighbours algorithms provided additional comparative baselines. Neural Network models were additionally used for nonlinear modelling, while Explainable AI techniques were applied to support feature interpretation.

Hyperparameter tuning was performed for neural network models, including hidden layer configuration, number of neurons, batch size, activation functions, optimizers, and training epochs. The experimental workflow also included comparative analysis of preprocessing strategies and feature selection approaches across different model configurations.

3.3. Evaluation strategy

Model performance was evaluated using several classification quality metrics, including Accuracy (ACC), Precision, Recall, F1-score, and ROC AUC. These metrics were selected in order to provide a comprehensive evaluation of model performance under class imbalance conditions.

To reduce dependence on random train-test splits and improve the reliability of performance estimates, both 5-fold and 10-fold cross-validation strategies were applied during the experimental evaluation process.

In addition to predictive performance analysis, model interpretability was evaluated using SHAP-based feature importance analysis. This approach enabled the identification of variables contributing most strongly to model predictions and supported the interpretation of factors associated with student dropout risk.

4. Results and discussion

4.1. Dataset comparison

Three experimental datasets (DR_01–DR_03) were constructed to evaluate the impact of different data preprocessing strategies on student dropout prediction performance. DR_01 included only the *Dropout* and *Graduate* categories, resulting in a binary classification problem with 3,630 records. DR_02 represented a logarithmically transformed version of DR_01, while DR_03 retained the original dataset structure and included observations from the *Enrolled* category that were reassigned using a neural network model.

Different preprocessing strategies were evaluated, including logarithmic transformation, normalization, and class balancing. These procedures generally improved model stability and classification performance, particularly for algorithms sensitive to feature scaling, such as SVM, k-NN, and neural networks. The comparative analysis demonstrated that classification performance depended not only on the selected machine learning algorithm but also on the dataset configuration and preprocessing approach. Table 2 compares the best-performing models across the three experimental datasets.

The highest classification performance was achieved on DR_02, where Logistic Regression obtained an accuracy of 0.909 and an F1-score of 0.927. However, the improvement compared with DR_01 was marginal. The logarithmic transformation slightly improved the performance of several models, particularly Decision Tree and Random Forest, while maintaining comparable results for Logistic Regression and Support Vector Machine. In contrast, DR_03 produced noticeably lower ACC and F1-score values, indicating that the inclusion of the processed *Enrolled* category increased classification complexity and reduced overall predictive performance.

Random Forest demonstrated strong performance across all experimental datasets. One possible explanation is its ability to capture complex relationships between academic, financial, and demographic variables without requiring extensive data transformation. Logistic Regression also achieved high accuracy despite its simpler structure, indicating that the most important factors were strongly associated with student dropout status. These findings suggest that both linear and tree-based approaches are suitable for student dropout prediction when relevant academic and financial indicators are available.

Table 2. Performance comparison of DR_01, DR_02 and DR_03 datasets

Dataset	Best model	ACC	F1-score	ROC AUC
DR_01	LR	0.905	0.924	0.96
DR_02	LR	0.909	0.927	0.96
DR_03	RF	0.858	0.899	0.89

The obtained results are consistent with previous studies reporting strong performance of Logistic Regression and Random Forest models in student dropout prediction. Similar findings were reported by Aulck et al. (2017), Naseem et al. (2019), and Miguéis et al. (2018), who also identified academic performance indicators as important predictors of student success and dropout risk. The accuracy values achieved in the present study are comparable to or higher than those reported in several earlier investigations, confirming the suitability of these approaches for educational data mining applications.

Although DR_02 achieved slightly higher evaluation metrics, DR_01 was selected for further experiments. The performance difference between the two datasets was minimal, while DR_01 preserved the original feature representation and avoided additional transformations. Furthermore, excluding the Enrolled category simplified the classification task and provided a clearer separation between the Dropout and Graduate classes. In addition, DR_01 demonstrated stable cross-validation results and lower sensitivity to parameter selection. These characteristics make DR_01 the most suitable dataset configuration for the subsequent feature selection and model evaluation experiments.

4.2. Feature selection

To reduce feature space redundancy and improve model interpretability, feature importance analysis was performed using correlation analysis and the SHAP (SHapley Additive Explanations) method. Initially, highly correlated variables were identified and removed from the dataset. The application of correlation-based filtering reduced data redundancy and computational complexity while preserving classification performance, indicating that several attributes contained overlapping information.

To further improve interpretability, the SHAP method was applied to identify the most influential features contributing to student dropout prediction. Unlike traditional feature importance approaches, SHAP provides both global and local explanations of model behaviour, enabling a more detailed analysis of feature contributions. Figure 2 shows the SHAP-based feature importance results.

The SHAP analysis revealed that the most influential variables were related to academic performance during the first and second semesters. In particular, the number of approved curricular units, semester grades, tuition fee status, and study programme demonstrated the strongest impact on prediction outcomes. Financial indicators, including tuition fee payments, also showed substantial influence on dropout risk. Table 3 presents the classification results obtained using the SHAP-selected features.

The obtained results demonstrate that feature selection based on SHAP preserved high classification performance while reducing the dimensionality of the dataset. Random Forest achieved the highest overall performance (ACC = 0.891, F1-score = 0.914, ROC AUC = 0.94), closely followed by Support Vector Machine. The findings indicate that a relatively small subset of informative features is sufficient to achieve reliable dropout prediction results.

The identified factors confirm that academic achievement and financial stability are the most important determinants of student dropout. Furthermore, the application of Explainable Artificial Intelligence methods improved model transparency and provided valuable insights into the factors associated with academic risk.

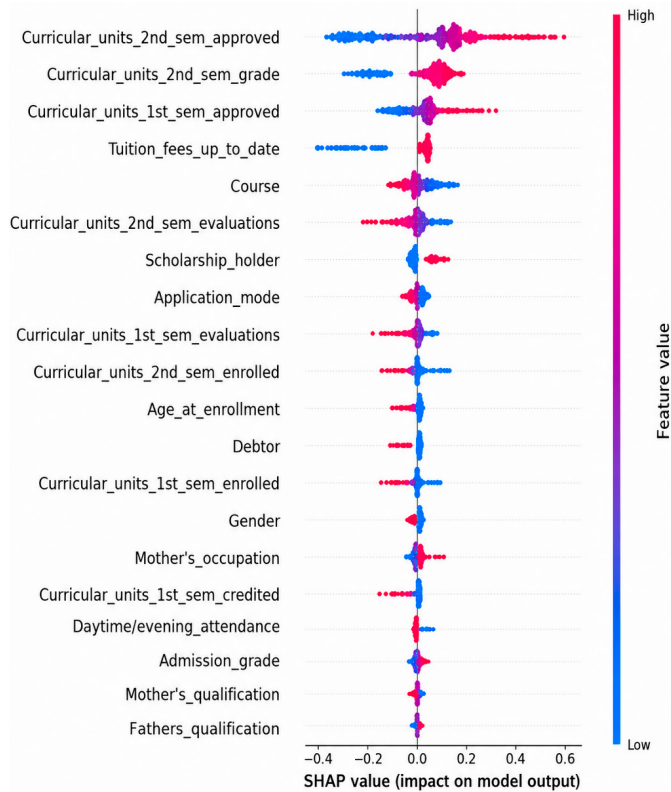


Figure 2. SHAP summary plot of feature importance

Table 3. Performance results using SHAP-selected features

Model	ACC	Precision	Recall	F1-score	ROC AUC
SVM	0.888	0.890	0.935	0.912	0.93
LR	0.870	0.874	0.924	0.898	0.90
DT	0.865	0.882	0.902	0.892	0.87
k-NN	0.870	0.872	0.927	0.900	0.88
RF	0.891	0.892	0.938	0.914	0.94

4.3. Model comparison

The comparative analysis was performed using several feature subsets to evaluate the impact of academic and financial information on student dropout prediction performance. The investigated configurations included financial and academic data, financial and first-semester academic data, pre-enrolment information only, and features selected using regression and correlation analysis. Table 4 presents the performance results for the different feature subsets.

Table 4. Performance comparison using different feature subsets

Feature subset	Best model	ACC	F1-score	ROC AUC
Financial and academic data	SVM/LR	0.902	0.922	0.95
Financial and first-semester data	RF	0.882	0.908	0.94
Pre-enrolment data only	RF	0.696	0.761	0.74
Regression and correlation selected features	LR	0.906	0.925	0.94

The results indicate that academic information had the strongest influence on classification performance. The highest accuracy (ACC = 0.906) and F1-score (0.925) were achieved using the feature subset selected through regression and correlation analysis. Similarly, the combination of financial and academic variables produced high predictive performance across multiple machine learning algorithms.

In contrast, the use of pre-enrolment information alone resulted in substantially lower performance, with the best model achieving an accuracy of only 0.696. This finding suggests that demographic and admission-related variables are insufficient for reliable dropout prediction when academic performance indicators are unavailable. Figure 3 shows the accuracy comparison across the different experimental feature subset configurations.

The graphical comparison confirms that models incorporating academic achievement indicators consistently outperformed those based solely on pre-enrolment characteristics. The results further demonstrate that a reduced subset of carefully selected features can achieve performance comparable to that obtained using the full dataset.

Overall, Logistic Regression, Support Vector Machine, and Random Forest demonstrated the most stable and accurate results across different experimental configurations, while Decision Tree and k-NN generally achieved lower predictive performance.

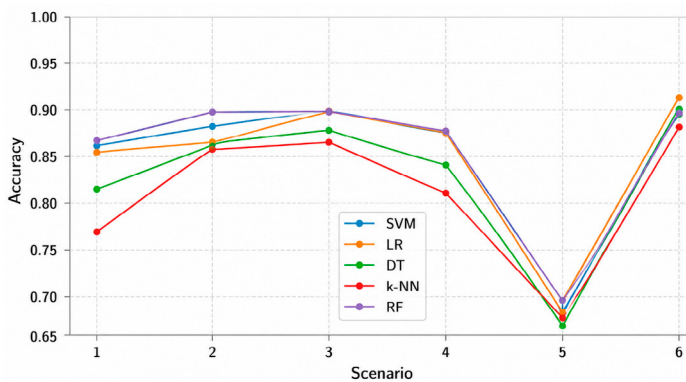


Figure 3. Accuracy comparison across six experimental feature subset configurations. Configurations 1–6 represent different combinations of academic, financial, pre-enrolment, and feature-selected datasets evaluated during the study

4.4. Important risk factors

The analysis identified several factors that have the greatest influence on student dropout risk. The most important predictors were associated with academic performance during the first and second semesters. In particular, the number of successfully completed curricular units, semester grades, and course completion indicators demonstrated the strongest contribution to model predictions. These findings are consistent with previous Educational Data Mining studies, which identify early academic performance as one of the most reliable indicators of student success and retention.

Financial variables also played an important role in dropout prediction. Students with outstanding tuition payments or overdue fees exhibited a substantially higher probability of dropout, whereas scholarship recipients were more likely to complete their studies successfully. These results highlight the importance of financial stability in supporting academic achievement and reducing dropout risk.

The SHAP analysis provided additional insight into the practical significance of these factors. Academic indicators, such as the number of approved courses and semester grades, can be monitored throughout the academic year to identify students who may require additional support. Financial variables, including tuition payment status and scholarship availability, may serve as early warning indicators and help institutions target financial assistance more effectively. The results also demonstrate that interpretable machine learning approaches can support decision-making by providing transparent explanations of individual predictions.

The analysis further revealed the influence of age at enrolment and study programme. Older students generally demonstrated a higher risk of academic difficulties, which may be associated with balancing studies, employment, and family responsibilities. Differences between study programmes suggest that programme-specific characteristics and academic workload may also affect student retention.

The identified risk factors have important practical implications for higher education institutions. By monitoring academic performance and financial indicators during the early stages of study, universities can identify at-risk students and implement targeted intervention measures. Such information can support the development of Early Warning Systems aimed at improving student retention and academic success.

The use of Explainable Artificial Intelligence techniques further improves the transparency of prediction models, enabling higher education institutions to better understand the factors contributing to student dropout risk.

4.5. Practical implications

The practical application of the proposed approach may include automated monitoring of students' academic performance, early identification of students at high risk of dropout, personalized academic advising, provision of additional financial support, and adaptation of student support strategies based on individual risk profiles.

The identified risk factors may also support resource allocation and intervention planning within higher education institutions. For example, students demonstrating declining academic performance or experiencing financial difficulties may be prioritized for academic

counselling, mentoring programmes, or financial assistance. Such targeted interventions may help institutions use available support resources more efficiently and improve student retention outcomes.

Particular importance is associated with the application of Explainable Artificial Intelligence methods in educational analytics. The use of SHAP analysis improves the transparency of machine learning models and enables explanation of the factors contributing to predictions for individual students. This is especially important in the practical implementation of intelligent decision-support systems within higher education institutions, where interpretability and algorithm transparency represent essential requirements.

Thus, the obtained results confirm the strong potential of machine learning and Explainable AI methods for educational analytics and academic risk prediction. The proposed approach may serve as a foundation for the development of intelligent decision-support systems aimed at improving student retention and enhancing the overall quality of the educational process.

5. Conclusions

This study examined the use of machine learning methods for student dropout prediction using academic, financial, and demographic data. Three dataset configurations and several feature selection approaches were evaluated to identify the most suitable combination for dropout prediction.

The results showed that dataset structure and feature selection had a noticeable impact on model performance. Among the investigated datasets, DR_01 was selected for further analysis because it provided stable results and a clear separation between the Dropout and Graduate classes. Correlation analysis and SHAP helped reduce the number of variables while maintaining high prediction accuracy.

The most important factors related to student dropout were academic indicators from the first and second semesters, particularly the number of approved curricular units and semester grades. Financial variables, including tuition fee status and scholarship support, also played an important role. These findings indicate that both academic performance and financial stability are closely related to student retention.

Logistic Regression, Random Forest, and Support Vector Machine achieved the best results across the evaluated experiments. The highest prediction performance was obtained when academic and financial variables were used together, while models based only on pre-enrolment information produced considerably lower accuracy.

The findings show that machine learning methods can be effectively applied to student dropout prediction and may support the development of Early Warning Systems in higher education institutions. In addition, the use of SHAP improved the interpretability of the models and provided a clearer understanding of the factors associated with dropout risk.

This study has several limitations. The experiments were performed using a single publicly available dataset, which may restrict the generalizability of the results. In addition, the analysed dataset represents a specific educational context and may not fully reflect the characteristics of students from other institutions, countries, or study systems. Future studies should evaluate the proposed approach on institutional datasets from different higher education environments and investigate the integration of additional behavioural and learning analytics variables.

Conflict of interests

The author declares no conflict of interest.

Author contributions

Conceptualization, methodology, software, validation, formal analysis, investigation, writing—original draft preparation, writing—review and editing, and visualization: A. Mincevičius.

References

- Aulck, L., Velagapudi, N., Blumenstock, J., & West, J. (2017). *Predicting student dropout in higher education*. ArXiv. <https://doi.org/10.48550/arXiv.1606.06364>
- Astin, A. W. (1993). *What matters in college? Four critical years revisited*. Jossey-Bass.
- Baker, R. S. J. d., & Yacef, K. (2009). The state of educational data mining in 2009: A review and future visions. *Journal of Educational Data Mining*, 1(1), 3–17. <https://doi.org/10.5281/zenodo.3554657>
- Bean, J. P. (1980). Dropouts and turnover: The synthesis and test of a causal model of student attrition. *Research in Higher Education*, 12(2), 155–187. <https://doi.org/10.1007/BF00976194>
- Behr, A., Giese, M., Teguim Kamdjou, H. D., & Theune, K. (2021). Dropping out of university: A literature review. *Review of Education*, 9(2), Article e3298. <https://doi.org/10.1002/rev3.3298>
- Cabrera, A. F., Nora, A., & Castaneda, M. B. (1992). The role of finances in the persistence process: A structural model. *Research in Higher Education*, 33(5), 571–593. <https://doi.org/10.1007/BF00973759>
- Carballo-Mendivil, B., Arellano-González, A., Ríos-Vázquez, N. J., & Lizardi-Duarte, M. d. P. (2025). Predicting student dropout from day one: XGBoost-based early warning system using pre-enrollment data. *Applied Sciences*, 15(16), Article 9202. <https://doi.org/10.3390/app15169202>
- Del Bonifro, F., Gabbrielli, M., Lisanti, G., & Zingaro, S. P. (2020). Student dropout prediction. In I. Bittencourt, M. Cukurova, K. Muldner, R. Luckin, & E. Millán (Eds.), *Lecture notes in computer science: Vol. 12163. Artificial Intelligence in education. AIED 2020* (pp. 129–140). Springer. https://doi.org/10.1007/978-3-030-52237-7_11
- Gallego, M. G., Perez de los Cobos, A. P., & Gallego, J. C. G. (2021). Identifying students at risk to academic dropout in higher education. *Education Sciences*, 11(8), Article 427. <https://doi.org/10.3390/educsci11080427>
- Goodfellow, I., Bengio, Y., & Courville, A. (2016). *Deep Learning*. MIT Press.
- Jiang, S., Williams, A. E., Schenke, K., Warschauer, M., & O'Dowd, D. (2014). Predicting MOOC performance with week 1 behavior. In *Proceedings of the 7th International Conference on Educational Data Mining (EDM 2014)* (pp. 273–275). https://educationaldatamining.org/EDM2014/uploads/procs2014/short%20papers/273_EDM-2014-Short.pdf
- Kehm, B. M., Larsen, M. R., & Sommersel, H. B. (2020). Student dropout from universities in Europe: A review of empirical literature. *Hungarian Educational Research Journal*, 10(2), 147–164. <https://doi.org/10.1556/063.2020.00018>
- Khan, W. Z., & Zubaidy, S. Al. (2017). Prediction of student performance in academic and military learning environment: Use of multiple linear regression predictive model and hypothesis testing. *International Journal of Higher Education*, 6(4), 152–160. <https://doi.org/10.5430/ijhe.v6n4p152>
- Larsen, M. R., Sommersel, H. B., & Larsen, M. S. (2013). *Evidence on dropout phenomena at universities*. Danish Clearinghouse for Educational Research.
- Lundberg, S. M., & Lee, S.-I. (2017). *A unified approach to interpreting model predictions*. ArXiv. <https://doi.org/10.48550/arXiv.1705.07874>
- Marcolino, M. R., Porto, T. R., Primo, T. T., Targino, R., Ramos, V., Queiroga, E. M., Munoz, R., & Cechinel, C. (2025). Student dropout prediction through machine learning optimization: Insights from Moodle log data. *Scientific Reports*, 15, Article 9840. <https://doi.org/10.1038/s41598-025-93918-1>

- Miguéis, V. L., Freitas, A., Garcia, P. J. V., & Silva, A. (2018). Early segmentation of students according to their academic performance: A predictive modelling approach. *Decision Support Systems*, 115, 36–51. <https://doi.org/10.1016/j.dss.2018.09.001>
- Naseem, M., Chaudhary, K., Sharma, B., & Lal, A. G. (2019). Using ensemble decision tree model to predict student dropout in computing science. In *2019 IEEE Asia-Pacific Conference on Computer Science and Data Engineering, CSDE 2019*. IEEE. <https://doi.org/10.1109/CSDE48274.2019.9162389>
- Orozco, C. M. E., & Niguidula, J. D. (2018). Student retention prediction using data mining techniques. *International Journal of Advanced Computer Science and Engineering*, 7(4), 9–14.
- Pascarella, E. T., & Terenzini, P. T. (2005). *How college affects students: A third decade of research* (Vol. 2). Jossey-Bass.
- Romero, C., & Ventura, S. (2010). Educational data mining: A review of the state of the art. *IEEE Transactions on Systems, Man, and Cybernetics, Part C (Applications and Reviews)*, 40(6), 601–618. <https://doi.org/10.1109/TSMCC.2010.2053532>
- Slade, S., & Prinsloo, P. (2013). Learning analytics: Ethical issues and dilemmas. *American Behavioral Scientist*, 57(10), 1510–1529. <https://doi.org/10.1177/0002764213479366>
- Spady, W. G. (1971). Dropouts from higher education: An interdisciplinary review and synthesis. *Interchange*, 1(1), 64–85. <https://doi.org/10.1007/BF02214313>
- Stinebrickner, R., & Stinebrickner, T. (2008). The effect of credit constraints on the college dropout decision: A direct approach using a new panel study. *American Economic Review*, 98(5), 2163–2184. <https://doi.org/10.1257/aer.98.5.2163>
- Suhirman, X. X., Zain, J. M., & Herawan, T. (2014). Data mining for education decision support: A review. *International Journal of Emerging Technologies in Learning*, 9(6), 4–19. <https://doi.org/10.3991/ijet.v9i6.3950>
- Terenzini, P. T., Rendón, L. I., Upcraft, M. L., Millar, S. B., Allison, K. W., Gregg, P. L., & Jalomo, R. (1994). The transition to college: Diverse students, diverse stories. *Research in Higher Education*, 35(1), 57–73. <https://doi.org/10.1007/BF02496662>
- Tinto, V. (1975). Dropout from higher education: A theoretical synthesis of recent research. *Review of Educational Research*, 45(1), 89–125. <https://doi.org/10.3102/00346543045001089>
- Tinto, V. (1987). *Leaving college: Rethinking the causes and cures of student attrition*. University of Chicago Press.
- Zhao, H., Zarar, S., Tashev, I., & Lee, C.-H. (2018). Convolutional-recurrent neural networks for speech enhancement. In *2018 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*. IEEE. <https://doi.org/10.1109/ICASSP.2018.8462155>