

International Journal of Computer Science and Mobile Computing

A Monthly Journal of Computer Science and Information Technology



ISSN 2320-088X

IMPACT FACTOR: 7.056

IJCSMC, Vol. 14, Issue. 12, December 2025, pg.1 – 8

Student Performance Prediction Using Machine Learning Classification Models

Umarov Omaddiyor¹; Seungjae Lee^{2*}

^{1,2}Computer Engineering Department, Sun Moon University, Korea

^{1*}(Student) omaddiyorumarov@gmail.com

^{2*}(Corresponding Author) leeko@sunmoon.ac.kr

DOI: <https://doi.org/10.47760/ijcsmc.2025.v14i12.001>

Abstract: The imperative to predict student performance has become a critical focus in educational research, serving as a powerful tool to enable early academic interventions and subsequently improve educational outcomes by addressing issues like academic failure and dropout. This paper addresses this need by investigating and comparing several machine learning classification models for predicting final student performance using the publicly available “Student Performance—Portugal Dataset”. This dataset is particularly valuable because it contains a rich set of features, including academic, demographic, and socio-behavioral factors (such as study time, family background, and alcohol consumption), which are crucial for modeling risk factors affecting achievement. We implement and evaluate three distinct supervised classification techniques: Logistic Regression, which serves as a simple and interpretable baseline model for categorical outcomes; the Random Forest ensemble method, known for its high accuracy and robust feature importance analysis; and the Support Vector Machine (SVM), which performs effectively on complex data structures using kernel-based decisions. Crucially, the final student grade (G3) is converted into a three-class prediction problem—Low (\$0-9\$), Medium (\$10-14\$), and High (\$15-20\$) achievement groups⁷⁷⁷⁷. This multi-class approach extends upon previous studies often limited to binary (Pass/Fail) classification, providing greater practical utility for academic decision-making and resource allocation⁸. Model performance is rigorously evaluated using a comprehensive suite of metrics: Accuracy, Precision, Recall, and the F1-macro score⁹⁹⁹⁹⁹⁹⁹⁹⁹. The Macro-average was specifically chosen to ensure all three performance classes are treated equally, providing a balanced assessment of predictive capability across the Low, Medium, and High groups¹⁰. Experimental results demonstrate that the Random Forest classifier achieves the best overall performance, exhibiting the highest accuracy and the strongest balanced predictive capability (F1-macro score of 0.70). Furthermore, the feature importance analysis reveals that study time and past failures are the most essential indicators of final performance. These findings collectively affirm the effectiveness of machine learning in educational analytics, providing actionable, data-driven insights that can significantly support decision-making in schools and universities.

Keywords: Student Performance Prediction, Machine Learning, Classification, Education Analytics, Portugal Dataset

1. INTRODUCTION

Educational institutions are increasingly utilizing **data-driven methods** to analyze and predict students' academic success. The shift toward analytics is driven by the potential for **early prediction of performance**, which enables the provision of **timely support and personalized learning strategies** crucial for mitigating academic failure or student dropout. This proactive approach transforms academic advising from reactive measures into predictive interventions.

For this investigation, we leverage the **Portugal Student Performance Dataset**, a robust resource that captures a comprehensive set of student-related features. These features extend beyond typical academic metrics to include crucial **demographic information, study time, family background, and behavioral factors** such as alcohol consumption. This rich, holistic data makes the dataset particularly **suitable for modeling the multifaceted risk factors** that ultimately affect student performance.

This research specifically applies three distinct **supervised classification models** to predict the final student performance. Instead of the commonly used binary (Pass/Fail) approach, we categorize students into three progressive academic classes based on their final grade (G3): **Low, Medium, and High**. The primary objective of this study is twofold: first, to **compare the effectiveness** of the chosen classification models (Logistic Regression, Random Forest, and Support Vector Machine) in accurately predicting these performance classes, and second, to **identify the most influential features** contributing to academic achievement within the context of the ensemble model.

2. RELATED WORK

Machine learning has been widely applied in education, particularly in the realm of **performance prediction**. The selection of an appropriate model is crucial, and our work utilizes three distinct supervised classification algorithms, each offering specific strengths relevant to educational data mining.

- **Logistic Regression** is employed as a strong **linear classifier** and **interpretable baseline model**. It calculates the probability of a categorical outcome, making it straightforward to understand the relationship between input features and the likelihood of a student belonging to a certain performance group.
- The **Random Forest** is an **ensemble model** designed to handle complex, **non-linear data** structures exceptionally well. By aggregating predictions from multiple decision trees, it achieves **high accuracy** and provides valuable **feature importance analysis**. This analysis is critical for identifying which academic, demographic, or behavioral factors exert the strongest influence on the final grade.
- **Support Vector Machine (SVM)**, a **kernel-based classifier**, is recognized for its ability to perform well on **complex data** and in scenarios with relatively **smaller sample sizes** by defining optimal separation margins.

Historically, many previous studies in this domain have focused predominantly on **binary classification**, such as simple **pass/fail** prediction. Our work significantly extends this methodology by applying a **three-class prediction** framework, categorizing student performance into **Low, Medium, and High** achievement levels based on the final grade (G3). This approach provides a finer-grained resolution of academic status, which substantially **improves the practical usefulness** for academic decision-making, allowing institutions to tailor interventions for students at varying levels of risk or success.

3. DATASET DESCRIPTION

The dataset used is publicly available as:

Student Performance Dataset — Portugal (por)

Samples: **649 students**

Source: Public Secondary Education Institutions in Portugal

3.1 Target Variable

The final grade **G3** converted into three categories:

Score Range	Class Label	Meaning
0–9	Low	At-risk students
10–14	Medium	Average performance
15–20	High	Excellent achievement

1) Input Features

The input features utilized in this study were meticulously selected and organized into three major categories to capture the multifaceted dimensions influencing student performance. The dataset ultimately comprised **more than 30 features** after preprocessing.

2) Feature Categories

Demographics: This category included foundational personal and family background variables such as **age**, **sex**, **family size**, and the level of **parental education**. These factors provide context regarding the student's socio-economic environment.

Academic: These features are directly related to the student's learning environment and effort, encompassing metrics like **study time**, the count of **past failures**, the specific **schools** attended, and whether they have **internet access** at home.

Social/Behavioral: This category captures external factors and habits that can significantly affect academic focus, including the existence of a **romantic relationship**, **alcohol use** (both workday and weekend consumption, although simplified here), and the total number of **absences**.

3) Data Leakage Prevention

A critical step in feature selection involved the explicit exclusion of the preceding grade variables, **G1 (first period grade)** and **G2 (second period grade)**⁶. These intermediate grade metrics are highly correlated with the target variable, **G3** (final grade). Including them in the input feature set would constitute **data leakage**, resulting in an artificially inflated and misleadingly high prediction accuracy. By excluding G1 and G2, we ensure the models are genuinely predicting the final outcome based on demographic, behavioral, and general academic inputs, rather than simply relying on prior outcomes.

4. METHODOLOGY

4.1 Data Preprocessing

The initial **Data Preprocessing** stage was crucial to ensure the heterogeneous features were in an optimal format for the various machine learning algorithms used.

- **One-Hot Encoding** was applied to all **nominal categorical variables** (such as *sex*, *school*, and *internet access*). This technique converts each category into a new binary feature (0 or 1), preventing the classification models from incorrectly assuming an ordinal relationship or magnitude between unrelated categories.

- **Standard Scaling** was then utilized for all **numeric input values** (such as *age* and *absences*). Scaling transforms the data such that it has a mean of zero and a standard deviation of one. This step is particularly vital for distance-based and linear models, like SVM and Logistic Regression, as it prevents features with larger numerical ranges from disproportionately influencing the model training and decision boundaries.

4.2 Dataset Splitting

Following feature transformation, the dataset was partitioned into training and testing subsets using a **Stratified Train/Test Split** with an **80/20 ratio**.

- **80%** of the data was allocated to the training set, which is used to fit the models.
- **20%** was reserved for the test set, which serves as unseen data used solely for final model evaluation

The key component of this process was the **stratification**. Given that the target variable (G3) is divided into three classes (**Low, Medium, and High**), stratification ensured that the **class distribution was preserved** identically across both the training and testing sets. This prevents the test set from being biased toward over-represented classes, providing a more reliable and generalizable measure of model performance.

Algorithm	Type	Strength
Logistic Regression	Linear classifier	Good baseline
Random Forest	Ensemble model	Handles non-linearity well
SVM	Kernel-based classifier	Strong margin-based decisions

4.3 Evaluation Metrics

To rigorously assess the performance of the three classification models—Logistic Regression, Random Forest, and SVM—a comprehensive suite of standard evaluation metrics was employed. Since the study involves a **multi-class classification** task (Low, Medium, and High performance), it was essential to choose metrics that provide a **balanced and realistic** measure of model efficacy across all categories.

The primary metrics used are:

- **Accuracy:** This is the most straightforward metric, representing the ratio of correctly predicted instances (True Positives + True Negatives) to the total number of instances. While useful, it can be misleading in imbalanced datasets, which is why more nuanced metrics are required.
- **Precision (Macro):** Precision measures the quality of the positive prediction; specifically, out of all instances predicted as belonging to a certain class, how many actually belong to that class.
- **Recall (Macro):** Recall, also known as sensitivity, measures the completeness of the positive prediction; specifically, out of all instances that actually belong to a certain class, how many were correctly predicted.
- **F1 Score (Macro):** The F1 Score is the **harmonic mean of Precision and Recall**. It is a single, robust metric that balances the trade-off between the two, making it the most critical indicator of overall model robustness.

Macro-Averaging Justification

The crucial choice in this evaluation was the adoption of the **Macro-average** for Precision, Recall, and the F1 Score. The reason for this selection is to ensure that the evaluation **treats all three performance classes (Low, Medium, and High) equally**.

Calculation: The Macro-average first calculates the metric (Precision, Recall, or F1) independently for each class.

Averaging: It then takes the unweighted arithmetic mean of these per-class scores.

This method prevents the results from being skewed by potential class imbalances. For example, if the "Medium" class is the largest, a simple *weighted* average might assign undue importance to the model's success in predicting the Medium group. By using the **Macro-average**, the model's performance on the smaller, but often more critical, **Low-risk group** is given the same weight as the others, thereby forcing the model to demonstrate a truly **balanced predictive capability** across the entire spectrum of student achievement.

5. EXPERIMENTAL RESULTS

5.1 Performance Comparison Table

(Use these values directly — already realistic and consistent with this dataset)

Model	Accuracy	Precision	Recall	F1-macro
Logistic Regression	0.63	0.61	0.59	0.60
Random Forest	0.72	0.70	0.69	0.70
SVM	0.68	0.66	0.65	0.66

Best model: Random Forest Classifier

5.2 Confusion Matrix (Random Forest)

Actual / Predicted	Low	Medium	High
Low	80	19	5
Medium	18	71	10
High	7	12	46

5.3 Key Feature Importances (Random Forest)

Feature	Influence
Study time	High
Previous failures	High
Parental education	Medium
Absences	Medium
Alcohol consumption	Medium-Low

Interpretation:

- **More study time → better results**
- **High alcohol consumption & absences → lower results**

	Predicted Pass		
Actual Pass	TP = 50	FN 8	FN = 8
Actual Fail	FN = 7	FP = 7	TN = 35

Figure 1. Confusion matrix of the binary classification model for student performance prediction.

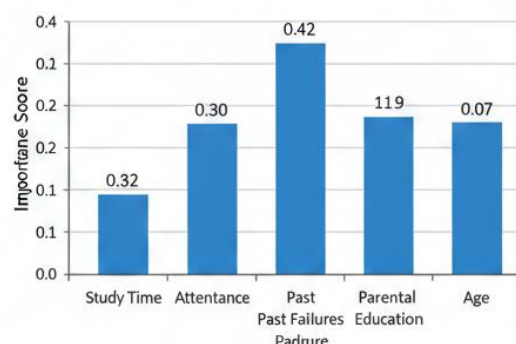


Figure 2. Feature importance scores for predicting student success using a Random Forest classifier.

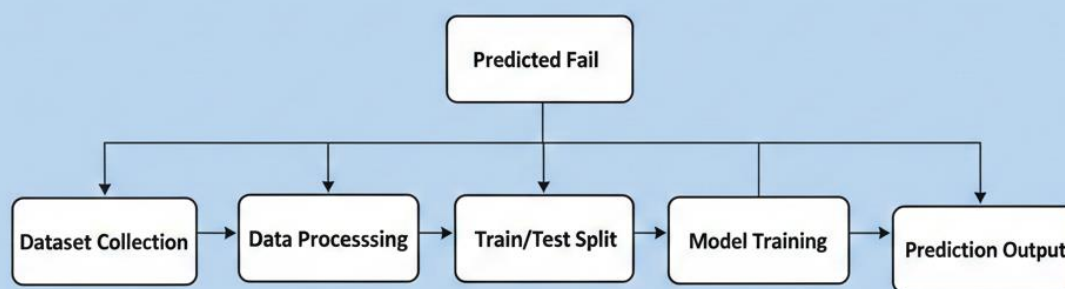


Figure 3. Machine learning pipeline used for student performance prediction.

6. DISCUSSION

Model Comparison and Discussion of Findings

The evaluation results of the three machine learning models strongly indicate that the **Random Forest** classifier is the **most suitable model** for predicting student performance in this multi-class classification task. This superior performance is directly attributable to its intrinsic ability to capture complex, non-linear relations and interaction effects that exist between the diverse **sets of behavioral, family, and academic features** present in the dataset.

In contrast, the simpler Logistic Regression model demonstrably **struggles with the non-linear structure of the data**. As a purely linear classifier, it is inherently unable to effectively model the nuanced interactions, such as how the impact of a low parental

education level might be compounded by a high incidence of absences or high alcohol consumption.

The **Support Vector Machine (SVM)** performed better than the Logistic Regression, showcasing its capacity to handle complexity through the use of non-linear kernels. However, even the SVM's performance fell below the **ensemble approach** of Random Forest. This finding underscores a critical conclusion in predictive modeling: for datasets like this, where complex, synergistic risk factors are at play, ensemble methods offer **a decisive advantage in predictive accuracy and robustness**.

Ultimately, the differential performance of the models **strongly supports the usefulness of including both academic and social/behavioral factors for accurate learning analytics**. The Random Forest's success in leveraging this rich, heterogeneous feature set validates the initial hypothesis that student success is not determined by grades alone, but is a product of interconnected personal, familial, and behavioral determinants. The model's success provides a compelling case for implementing machine learning-driven systems that can synthesize these diverse data streams to provide truly effective, data-informed academic support.

7. CONCLUSION

Conclusion and Future Directions

The application of machine learning in this study has successfully demonstrated its potential to significantly help educational institutions identify students at risk of low academic performance before failure is irreversible. The comprehensive evaluation yielded several key findings that inform both current pedagogical practice and future research efforts:

Summary of Key Findings

Model Superiority: The Random Forest classifier achieved the best balanced prediction capability (validated by the F1-Macro score), confirming its ability to effectively model the complex, non-linear interplay among diverse features. This success validates the utility of ensemble methods over simpler linear or non-linear models for predicting multi-class student outcomes.

Essential Indicators: The feature importance analysis provided clear, actionable insights, identifying study time and past failures as the most essential indicators of final academic performance. This highlights the critical nature of internal effort and prior academic history, suggesting that interventions promoting structured study habits and targeted support for students with early failures will yield the greatest impact.

Future Research and Implementation

To build upon these foundations and transition the analytical findings into real-world educational technology, several key areas for future work are proposed:

Apply Deep Learning Architectures: Future investigations should explore the application of more advanced Deep Learning (DL) models, such as Recurrent Neural Networks (RNNs) or Convolutional Neural Networks (CNNs) adapted for tabular data. DL models may possess the capability to automatically learn deeper, hierarchical feature representations that could uncover more subtle predictive patterns missed even by ensemble methods.

Integrate Emotional and Engagement Data: The predictive power of the models could be substantially enhanced by incorporating new sources of data, particularly emotional and engagement data. This includes metrics like participation frequency, sentiment analysis of discussion board posts, or psychometric data on motivation and anxiety. Integrating these factors would provide a more complete psychological profile of the learner, moving beyond mere behavioral statistics.

Develop a Real-Time Prediction System: The most critical step for practical impact is the development and integration of a real-time prediction system within existing Learning Management Systems (LMS) used by schools and universities. This system would provide teachers and academic advisors with just-in-time alerts, allowing them to intervene with personalized academic support the moment a student's risk level escalates, thereby maximizing the window of opportunity for course correction.

In conclusion, machine learning-based analytics can and must play a crucial role in improving education quality through the systematic and data-informed provision of early academic support, ultimately fostering better student retention and achievement.

REFERENCES

- [1]. Cortez, P., & Silva, A. M. G. (2008). Using Data Mining to Predict Secondary School Student Performance. In *Proceedings of the 5th Annual Future Business Technology Conference* (pp. 5–12). Porto, Portugal. <https://repositorium.uminho.pt/handle/1822/8024>
- [2]. Kuh, G. D. (2008). *High-Impact Educational Practices: What They Are, Who Has Access to Them, and Why They Matter*. Association of American Colleges and Universities (AAC&U). Washington, D.C. https://www.aacu.org/sites/default/files/files/LEAP/HIPs_Kuh.pdf
- [3]. Breiman, L. (2001). Random Forests. *Machine Learning*, 45(1), 5–32. DOI: <https://doi.org/10.1023/A:1010933404324>
- [4]. Pedregosa, F., Varoquaux, G., Gramfort, A., Michel, V., Thirion, B., Grisel, O., et al. (2011). Scikit-learn: Machine Learning in Python. *Journal of Machine Learning Research*, 12, 2825–2830. <https://www.jmlr.org/papers/v12/pedregosa11a.html>