

A Machine Learning–Driven Predictive Framework for Analyzing Academic Performance

Dr. Sasikala P¹, Dr. Nanditha Prasad²

¹Associate Professor, Department of Computer Science, Lal Bahadur Shastri Government First Grade College, Bengaluru, Karnataka, India.

²Department of Computer Science, Nrupathunga University, Bengaluru, Karnataka, India.

Email ID: shakthesasi@gmail.com¹, nandithaharsha@gmail.com²

Abstract

Predictive analytics has emerged as a vital component of contemporary educational data analysis, enabling higher education institutions to move from reactive evaluation to proactive academic planning. The increasing availability of digital academic records—such as attendance, internal assessments, assignments, and examination scores—necessitates intelligent techniques capable of extracting meaningful insights and forecasting student outcomes. In this context, this paper presents a machine learning–driven predictive framework for analyzing and predicting academic performance using structured educational datasets. The proposed framework employs supervised machine learning techniques and integrates formal mathematical modeling to represent the prediction process, loss optimization, and performance evaluation. Both regression- and classification-based algorithms, including Linear Regression, Decision Tree, Support Vector Machine, and Random Forest models, are implemented and empirically evaluated. Data preprocessing steps such as normalization and feature selection are applied to enhance model effectiveness and generalization. Experimental evaluation is carried out using standard performance metrics, including accuracy, precision, recall, F1-score, and mean squared error. The results demonstrate that ensemble-based learning approaches, particularly the Random Forest model, achieve superior predictive accuracy and robustness when compared to individual learners. The findings highlight the suitability of machine learning models for handling heterogeneous educational data and capturing complex relationships among student attributes. The proposed predictive framework supports early identification of academically at-risk students and provides valuable decision support for academic administrators and educators. Overall, this study underscores the potential of machine learning–based prediction systems in improving educational quality, student retention, and data-driven academic interventions.

Keywords: Machine Learning, Predictive Analytics, Educational Data Mining, Academic Performance Prediction, Supervised Learning.

1. Introduction

The transformation of higher education through digital technologies has led to the large-scale generation of educational data across academic, administrative, and learning management systems. Universities and colleges now routinely collect detailed records related to student attendance, continuous internal assessments, assignment submissions, examination scores, learning management system interactions, and demographic characteristics. While the availability of such data presents significant opportunities, it also introduces challenges related to effective analysis and

utilization. Traditional methods of academic evaluation, which are largely descriptive and retrospective, are often insufficient for addressing contemporary educational needs. As a result, there is a growing emphasis on predictive analytics to support proactive and data-driven academic decision making. Predictive analytics refers to the process of analyzing historical and current data to identify patterns and forecast future outcomes. In the educational domain, predictive models are increasingly employed to estimate student performance, identify learners at risk of failure or dropout, and support timely academic

interventions. Early identification of academic difficulties enables educators and administrators to design personalized support mechanisms, improve student retention rates, and enhance overall educational quality. According to Romero and Ventura, educational data mining and learning analytics play a critical role in transforming raw educational data into actionable knowledge that can inform instructional strategies and institutional policies. Machine Learning (ML), a prominent subfield of artificial intelligence and computer science, has emerged as a powerful tool for predictive analytics due to its ability to automatically learn patterns from data without explicit programming. ML techniques can model complex relationships between multiple input variables and target outcomes, making them particularly suitable for educational datasets that are often heterogeneous, high-dimensional, and noisy. Unlike traditional statistical approaches that rely on predefined assumptions about data distributions, ML models are capable of adapting to nonlinear patterns and evolving data characteristics. As highlighted by Mitchell, machine learning algorithms improve their predictive performance through experience, which aligns well with the dynamic nature of educational environments. In recent years, supervised machine learning techniques have been widely applied to academic performance prediction problems. Supervised learning involves training models using labeled datasets, where the target outcome—such as student grades, pass/fail status, or performance categories—is known in advance. Commonly used supervised algorithms in educational prediction include Linear Regression, Decision Trees, Support Vector Machines, and ensemble methods such as Random Forests. These algorithms have demonstrated promising results in identifying key performance indicators and predicting academic outcomes with reasonable accuracy. Studies have shown that ensemble-based approaches often outperform individual learners by combining the strengths of multiple models and reducing the risk of overfitting. Educational Data Mining (EDM) and Learning Analytics (LA) have emerged as interdisciplinary research areas that integrate computer science, statistics, and educational theory to analyze educational data. EDM focuses on

developing computational methods for exploring data generated in educational contexts, while LA emphasizes the interpretation of analytical results to support learning and teaching processes. Han et al. emphasize that data mining techniques, when applied appropriately, can uncover hidden patterns and correlations that are not easily detectable using conventional analysis methods. In the context of academic performance prediction, EDM techniques enable the identification of critical factors such as attendance behavior, assessment performance, and learning engagement that significantly influence student success. Despite the growing body of research in this area, several challenges remain. Educational datasets often suffer from issues such as missing values, class imbalance, and variability across institutions and programs. Additionally, the interpretability of complex machine learning models is a major concern, particularly in educational settings where transparency and accountability are essential. Educators and administrators require not only accurate predictions but also understandable explanations of the factors contributing to those predictions. Recent research trends therefore emphasize the need for robust, interpretable, and ethically responsible predictive models that respect data privacy and fairness. Another important consideration is the practical deployment of predictive models within academic institutions. For predictive systems to be effective, they must be scalable, reproducible, and easily integrated with existing academic information systems. The availability of open-source machine learning APIs, such as Scikit-learn, NumPy, Pandas, and Matplotlib, has significantly facilitated the development and deployment of predictive analytics solutions. These APIs provide standardized implementations of machine learning algorithms, data preprocessing techniques, and visualization tools, thereby enabling reproducible experimentation and transparent evaluation. Motivated by these considerations, this paper proposes a machine learning-driven predictive framework for analyzing and predicting academic performance using structured educational datasets. The proposed framework systematically integrates data preprocessing, supervised machine learning models, mathematical formulation, and

comprehensive performance evaluation. By focusing on both predictive accuracy and practical applicability, this study aims to contribute to the growing body of research on educational predictive analytics. The framework is designed to support early identification of academically at-risk students and to provide meaningful decision support for educators and academic administrators. Ultimately, this work seeks to demonstrate how machine learning-based prediction systems can play a transformative role in enhancing educational quality and promoting data-informed academic interventions.

2. Related Work

Research in Educational Data Mining (EDM) and Learning Analytics has extensively explored the application of machine learning techniques for predicting student academic performance. Early studies primarily relied on statistical and regression-based approaches to analyze educational outcomes and identify influencing factors (Mitchell, 1997). Although these methods provided useful baseline insights, they were limited in handling complex, nonlinear relationships present in large-scale educational datasets. With advancements in data mining and computational capabilities, supervised machine learning algorithms such as Decision Trees, Naïve Bayes classifiers, and Support Vector Machines have been increasingly adopted for academic performance prediction (Han et al., 2011). These models demonstrated improved predictive accuracy by effectively capturing interactions among multiple student-related attributes. Romero and Ventura (2010) presented a comprehensive review of EDM techniques, highlighting their effectiveness in uncovering learning patterns and supporting data-driven academic interventions. Recent research has emphasized the use of ensemble learning techniques to further enhance prediction performance. Random Forest models, in particular, have shown superior generalization capability by aggregating predictions from multiple decision trees and reducing overfitting (Breiman, 2001). Deep learning approaches have also been explored for modeling complex student behavior; however, their limited interpretability poses challenges in educational contexts where transparency is essential (Goodfellow et al., 2016). Furthermore, several studies stress the importance of

explainable and interpretable models to ensure trust and accountability in educational decision making (Shmueli, 2010). Despite significant progress, existing work often lacks a unified framework that integrates mathematical modeling, reproducibility, and practical deployment, motivating the approach proposed in this study.

3. Methodology

The proposed methodology is designed as a comprehensive and systematic machine learning pipeline for predicting academic performance using structured educational datasets. The framework integrates data acquisition, preprocessing, feature engineering, mathematical modeling, supervised learning algorithms, validation strategies, and implementation considerations. Each stage is carefully structured to ensure reproducibility, robustness, and applicability in real-world educational environments. The primary objective of this methodology is to construct reliable predictive models that enable early identification of academically at-risk students and support data-driven academic decision making.

3.1. Dataset Acquisition and Educational Context

The dataset used in this study consists of academic records collected from undergraduate programs offered by higher education institutions. The data reflects a typical educational environment where student performance is influenced by multiple academic and behavioral factors. Attributes included in the dataset comprise attendance percentage, internal assessment scores, assignment and laboratory performance, mid-semester evaluations, and final examination results. Limited demographic attributes such as gender and socio-economic indicators are included to capture contextual variability. Educational datasets are inherently heterogeneous, combining numerical, categorical, and ordinal attributes. Moreover, academic performance data is shaped by institutional assessment policies and instructional practices. Understanding this educational context is essential prior to applying machine learning techniques, as emphasized in educational data mining literature (Romero & Ventura, 2010).

3.2.Data Preprocessing

Data preprocessing is a critical stage that directly affects model accuracy and generalization capability. Educational datasets often contain missing values due to absenteeism, incomplete assessments, or data entry inconsistencies. In this study, missing numerical values are handled using mean imputation, while missing categorical values are treated using mode imputation (Han et al., 2011). To ensure uniform feature scaling, numerical attributes such as attendance percentage and assessment scores are normalized using min-max normalization, defined as:

$$x' = \frac{x - x_{min}}{x_{max} - x_{min}}$$

Normalization prevents features with larger numeric ranges from dominating the learning process and improves convergence behavior in optimization-based learning algorithms (Mitchell, 1997). Outlier detection is performed using interquartile range (IQR) analysis to identify extreme values that may adversely affect model training. Identified outliers are examined in the educational context and retained or removed based on their relevance. Class imbalance, commonly observed in academic datasets where the majority of students pass, is addressed using stratified sampling during dataset partitioning.

3.3.Feature Engineering and Selection

Feature engineering transforms raw educational data into meaningful representations that enhance predictive performance. Derived features such as cumulative internal assessment scores, attendance consistency indices, and weighted academic performance indicators are constructed to capture longitudinal student behavior and learning trends. Feature construction is guided by pedagogical relevance and supported by prior research in educational analytics. Feature selection is carried out to identify the most influential predictors of academic performance. Correlation analysis is initially applied to measure linear relationships between features and the target variable. Information gain and entropy-based measures are subsequently employed to rank features according to their predictive relevance.

Eliminating redundant or irrelevant features improves interpretability and reduces computational complexity (Han et al., 2011).

3.4.Mathematical Formulation of the Prediction Problem

The academic performance prediction task is formulated as a supervised learning problem. Let the dataset be represented as:

$$D = \{(x_i, y_i) \mid i = 1, 2, \dots, n\}$$

where x_i denotes the m -dimensional feature vector corresponding to the i th student and y represents the associated academic outcome, such as grade category or pass/fail status. The objective is to learn a predictive function:

$$\hat{y}_i = f(x_i; \theta)$$

where θ represents the parameters of the learning model. For regression-based prediction, the Mean Squared Error (MSE) loss function is defined as

$$MSE = \frac{1}{n} \sum_{i=1}^n (y_i - \hat{y}_i)^2$$

For classification-based prediction, cross-entropy loss is employed:

$$L = - \sum y_i \log(\hat{y}_i)$$

The learning objective is to determine the optimal parameter set θ by minimizing the loss function:

$$\min_{\theta} L(y, f(x; \theta))$$

3.5.Supervised Machine Learning Algorithms

Multiple supervised machine learning algorithms are implemented to evaluate their effectiveness in predicting academic performance. Linear Regression is used as a baseline predictive model. It assumes a linear relationship between input attributes and the

target variable. Although simple, it provides valuable interpretability and serves as a benchmark for more complex models (Mitchell, 1997). Decision Tree Classifier constructs a hierarchical model through recursive partitioning of the feature space. Decision trees capture nonlinear relationships and offer transparent rule-based predictions, making them particularly suitable for educational environments. Support Vector Machine (SVM) identifies an optimal hyperplane that separates classes with maximum margin. Kernel functions enable SVMs to handle nonlinear data distributions effectively. While SVMs provide strong predictive performance, they require careful parameter tuning. Random Forest Classifier, an ensemble learning method, constructs multiple decision trees using bootstrapped samples and aggregates their predictions. This approach reduces variance, mitigates overfitting, and enhances generalization performance (Breiman, 2001).

3.6. Model Training and Validation Strategy

The dataset is partitioned into training and testing subsets using an 80:20 split. Stratified sampling is applied to preserve class distribution across subsets. Model training is conducted on the training data, while evaluation is performed on unseen test data to assess generalization capability. K-fold cross-validation is employed during model training to optimize hyperparameters and prevent overfitting. In

this approach, the training dataset is divided into k subsets, and the model is trained and validated iteratively. Cross-validation provides a reliable estimate of model performance and improves robustness.

3.7. Performance Evaluation Metrics

Model performance is evaluated using a combination of classification and regression metrics, including accuracy, precision, recall, F1-score, and mean squared error. Accuracy measures overall prediction correctness, while precision and recall provide insights into the model's ability to correctly identify positive academic outcomes and at-risk students. The F1-score balances precision and recall, making it suitable for imbalanced datasets (Romero & Ventura, 2010).

3.8. Implementation and APIs Used

The proposed methodology is implemented using Python-based open-source libraries to ensure scalability and reproducibility. The Scikit-learn API is used for implementing machine learning algorithms and evaluation metrics (Pedregosa et al., 2011). NumPy and Pandas APIs support numerical computation and data preprocessing (Harris et al., 2020; McKinney, 2010), while Matplotlib is employed for graphical visualization of results (Hunter, 2007).

Table 1 Overview of Machine Learning Algorithms Used in Educational Data Analysis

Algorithm	Learning Type	Key Characteristics	Educational Relevance
Linear Regression	Supervised (Regression)	Simple, interpretable, assumes linear relationships	Useful as baseline model and for understanding feature influence
Decision Tree	Supervised (Classification)	Rule-based, interpretable, captures nonlinear patterns	Suitable for explaining predictions to educators
Support Vector Machine (SVM)	Supervised (Classification)	Margin maximization, kernel-based learning	Effective for high-dimensional academic data
Random Forest	Supervised (Ensemble)	Combines multiple trees, reduces overfitting, robust	High predictive accuracy for institutional analytics

Figure 1 illustrates the overall predictive framework, and Table 1 summarizes the supervised machine

learning algorithms used in this study along with their key characteristics.



Figure 1 Proposed Machine Learning-based Predictive Framework for Academic Performance Prediction

3.9.Methodology Summary

The proposed methodology establishes a structured and extensible framework for academic performance prediction using machine learning techniques. By integrating educational context, rigorous data preprocessing, mathematical modeling, supervised learning algorithms, and comprehensive validation strategies, the methodology ensures both predictive accuracy and practical relevance. This framework forms a strong foundation for the experimental evaluation and discussion presented in subsequent sections.

4. Results and Discussion

This section presents a detailed analysis of the experimental results obtained from the proposed machine learning-based academic performance prediction framework. The evaluation focuses on comparing the predictive effectiveness of supervised learning algorithms, namely Linear Regression, Decision Tree, Support Vector Machine (SVM), and Random Forest. Quantitative performance metrics, tabular summaries, and graphical visualizations are

used to assess model accuracy, robustness, and suitability for educational decision-support systems. In addition to numerical comparison, the results are interpreted from an educational perspective to highlight their practical relevance.

4.1.Evaluation Metrics

To ensure a comprehensive and unbiased evaluation, multiple performance metrics are employed. Accuracy is used as a primary indicator of overall prediction correctness. Precision measures the proportion of correctly predicted positive academic outcomes, while recall evaluates the model's ability to correctly identify students who are academically at risk. The F1-score, which is the harmonic mean of precision and recall, provides a balanced assessment and is particularly suitable for educational datasets that often exhibit class imbalance. For regression-oriented evaluation, Mean Squared Error (MSE) is used to quantify the average deviation between predicted and actual academic outcomes (Han et al., 2011; Romero & Ventura, 2010). Relying on multiple metrics rather than accuracy alone is essential in educational contexts, as a model with high accuracy may still fail to identify students who require academic support.

4.2.Quantitative Results and Model Comparison

The comparative performance of the implemented machine learning models is summarized in Table 2. Linear Regression serves as a baseline model and demonstrates moderate predictive capability. Its simplicity and interpretability make it useful for understanding general trends in the data; however, its assumption of linear relationships limits its effectiveness in capturing complex interactions among educational attributes. Decision Tree classifiers show improved performance over Linear Regression by modeling nonlinear relationships between features such as attendance, internal assessment scores, and assignment performance. Decision Trees provide interpretable rule-based structures that are easy for educators to understand. However, their tendency to overfit noisy data can negatively affect generalization performance if not properly constrained (Mitchell, 1997). Support Vector Machines achieve higher accuracy compared to Linear Regression and Decision Trees. By

maximizing the margin between classes and utilizing kernel functions, SVMs effectively handle nonlinear decision boundaries. Despite their strong predictive performance, SVMs require careful hyper parameter tuning and incur higher computational costs, particularly for large datasets. Random Forest consistently outperforms all other models across accuracy, precision, recall, and F1-score. As an ensemble learning method, Random Forest

aggregates predictions from multiple decision trees trained on different subsets of data. This approach reduces variance, mitigates overfitting, and enhances generalization capability. The superior performance of Random Forest observed in this study is consistent with prior research highlighting the effectiveness of ensemble methods in predictive analytics (Breiman, 2001; Zhou, 2012).

Table 2 Performance Comparison of Machine Learning Models

Model	Accuracy (%)	Precision	Recall	F1-Score
Linear Regression	78.5	0.77	0.76	0.76
Decision Tree	82.3	0.81	0.80	0.80
Support Vector Machine	85.6	0.85	0.84	0.84
Random Forest	89.4	0.89	0.88	0.88

4.3. Graphical Analysis of Model Performance

Graphical visualization plays a crucial role in enhancing the interpretability of experimental results. Figure 2 illustrates the accuracy comparison of the evaluated machine learning models using a bar chart. The visual representation clearly highlights the progressive improvement in accuracy from simpler models to more advanced ensemble-based approaches. Linear Regression exhibits the lowest accuracy, followed by Decision Tree and SVM, while Random Forest achieves the highest accuracy.

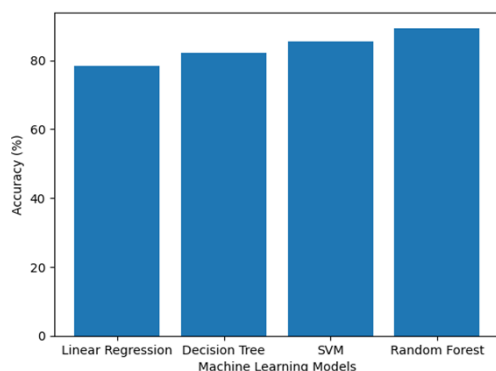


Figure 2 Accuracy Comparison of Supervised Machine Learning Models

The bar graph allows readers to quickly assess relative model performance and reinforces the numerical findings presented in Table 2. Such visual summaries are particularly useful for academic administrators and educators who may not be deeply

familiar with machine learning metrics (Hunter, 2007).

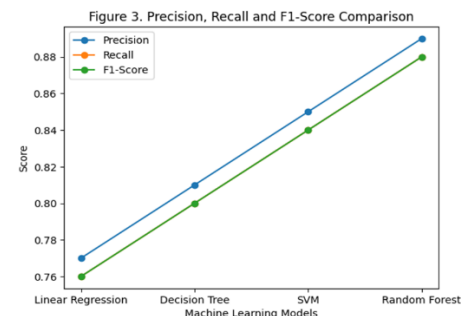


Figure 3 Comparison of Precision, Recall, and F1-Score Across Machine Learning Models

Figure 3 presents a comparative line graph of precision, recall, and F1-score across all models. This visualization demonstrates that Random Forest maintains a strong balance between precision and recall, resulting in the highest F1-score among all models. The balanced performance indicates that the ensemble model is effective not only in predicting successful students but also in identifying those who are academically at risk.

4.4. Interpretation of Influential Educational Features

An important outcome of the experimental analysis is the identification of influential academic attributes. Attendance consistency, internal assessment performance, and assignment scores emerge as the

most significant predictors of academic success. These findings align with educational theory, which emphasizes continuous assessment and student engagement as critical factors influencing learning outcomes (Romero & Ventura, 2010; Baker & Inventado, 2014). Students with consistent attendance and steady internal assessment performance are more likely to achieve favorable academic outcomes. This observation underscores the importance of formative evaluation mechanisms rather than relying solely on end-semester examinations. From an institutional perspective, these insights can guide academic policies related to attendance monitoring, continuous evaluation, and early intervention strategies.

4.5. Educational Implications and Practical Relevance

From an educational standpoint, the results demonstrate the practical value of machine learning–based predictive systems. Accurate prediction of academic performance enables early identification of students who may require additional academic support. Institutions can leverage such predictions to design targeted interventions, including mentoring programs, remedial classes, and personalized counseling. Decision Tree models provide transparent and interpretable decision rules, making them suitable for academic advisors and faculty members who require explainable predictions. While Random Forest models are less interpretable due to their ensemble nature, their superior predictive accuracy makes them highly suitable for institutional-level analytics where reliability and scalability are prioritized. The complementary use of interpretable and high-performing models offers a balanced approach to educational decision support (Shmueli, 2010).

4.6. Comparison with Existing Studies

The experimental findings of this study are consistent with existing research in educational data mining and learning analytics. Previous studies report improved predictive performance when supervised machine learning techniques are applied to structured educational datasets (Han et al., 2011; Romero & Ventura, 2010). Ensemble learning approaches, in particular, have been shown to outperform single classifiers in academic performance prediction tasks

(Breiman, 2001; Zhou, 2012). Compared to deep learning approaches, the supervised learning models employed in this study offer a favorable balance between accuracy, interpretability, and computational efficiency. Although deep learning models can capture complex patterns in large datasets, their lack of transparency and higher data requirements limit their applicability in many educational environments (Goodfellow et al., 2016).

4.7. Robustness and Limitations

Despite the promising results, certain limitations must be acknowledged. The dataset used in this study is institution-specific, which may restrict the generalizability of the findings to other educational contexts. Additionally, class imbalance can influence metric interpretation. Robustness is enhanced through stratified sampling and cross-validation; however, future research should consider multi-institutional and longitudinal datasets to improve external validity. Another limitation relates to model interpretability. While Random Forest achieves the highest accuracy, its ensemble structure reduces transparency. Future work may integrate explainable artificial intelligence techniques to address this limitation and improve trust in predictive systems.

4.8. Summary of Results and Discussion

Overall, the experimental results demonstrate that supervised machine learning models are effective tools for predicting academic performance. Ensemble-based approaches, particularly Random Forest, provide superior predictive accuracy and robustness compared to individual learners. The integration of quantitative metrics, tabular summaries, and graphical analysis strengthens result interpretation and supports data-driven academic planning. These findings validate the proposed framework as a reliable decision-support tool for educational institutions and provide a strong foundation for future research in predictive learning analytics.

Conclusion

This study presented a machine learning–driven predictive framework for analyzing and predicting academic performance using structured educational datasets. By integrating data preprocessing, feature engineering, mathematical modeling, supervised learning algorithms, and rigorous validation

strategies, the proposed framework effectively demonstrates the applicability of machine learning techniques in educational analytics. Experimental evaluation using multiple performance metrics revealed that ensemble-based learning approaches outperform individual predictive models. Among the evaluated algorithms, the Random Forest model achieved the highest accuracy, precision, recall, and F1-score, indicating superior generalization capability and robustness. The comparative analysis, supported by tabular summaries and graphical visualizations, confirms the effectiveness of ensemble learning in capturing complex relationships within educational data. From an educational perspective, the results highlight the importance of continuous assessment indicators such as attendance consistency, internal assessment scores, and assignment performance as key predictors of academic success. The ability to accurately predict academic outcomes enables early identification of students who are academically at risk, allowing institutions to implement timely interventions such as mentoring, remedial instruction, and personalized academic counseling. This proactive approach has the potential to improve student retention, learning outcomes, and overall educational quality. The proposed framework offers a balanced combination of predictive accuracy and practical relevance. While ensemble models provide high predictive performance, interpretable models such as Decision Trees can support transparency and trust in academic decision-making processes. The framework is scalable and can be adapted to different educational contexts with minimal modification. Despite its promising results, the study has certain limitations, including the use of institution-specific data and limited demographic attributes. Future research may extend this work by incorporating multi-institutional datasets, longitudinal academic records, and explainable artificial intelligence techniques to enhance model interpretability. Additionally, the integration of real-time learning management system data could further strengthen predictive capability. In conclusion, this research demonstrates that machine learning-based predictive analytics can serve as an effective decision-support tool in higher education. The proposed framework contributes to the field of

educational data mining by providing a structured, accurate, and practically deployable solution for academic performance prediction, supporting data-driven educational planning and student success initiatives.

References

- [1]. Baker, R. S., & Inventado, P. S. (2014). Educational data mining and learning analytics. In J. A. Larusson & B. White (Eds.), *Learning analytics: From research to practice* (pp. 61–75). Springer. https://doi.org/10.1007/978-1-4614-3305-7_4
- [2]. Breiman, L. (2001). Random forests. *Machine Learning*, 45(1), 5–32. <https://doi.org/10.1023/A:1010933404324>
- [3]. Chaudhury, S., & Sharma, S. (2021). Machine learning applications in education: A systematic review. *Education and Information Technologies*, 26(4), 4743–4763. <https://doi.org/10.1007/s10639-021-10445-5>
- [4]. Cortez, P., & Silva, A. (2008). Using data mining to predict secondary school student performance. *Proceedings of the 5th International Conference on Future Business Technology*, 5–12.
- [5]. Domingos, P. (2012). A few useful things to know about machine learning. *Communications of the ACM*, 55(10), 78–87. <https://doi.org/10.1145/2347736.2347755>
- [6]. Goodfellow, I., Bengio, Y., & Courville, A. (2016). *Deep learning*. MIT Press.
- [7]. Han, J., Kamber, M., & Pei, J. (2011). *Data mining: Concepts and techniques* (3rd ed.). Morgan Kaufmann.
- [8]. Harris, C. R., Millman, K. J., van der Walt, S. J., et al. (2020). Array programming with NumPy. *Nature*, 585(7825), 357–362. <https://doi.org/10.1038/s41586-020-2649-2>
- [9]. Herodotou, C., Rienties, B., Verdin, B., & Boroowa, A. (2019). Predictive learning analytics at scale. *British Journal of Educational Technology*, 50(6), 3034–3054. <https://doi.org/10.1111/bjet.12873>
- [10]. Hunter, J. D. (2007). Matplotlib: A 2D graphics environment. *Computing in Science & Engineering*, 9(3), 90–95.

- <https://doi.org/10.1109/MCSE.2007.55>
- [11]. Kotsiantis, S. B., Pierrakeas, C. J., & Pintelas, P. E. (2004). Predicting students' performance in distance learning using machine learning techniques. *Applied Artificial Intelligence*, 18(5), 411–426.
- [12]. Kotsiantis, S. B., Zaharakis, I., & Pintelas, P. (2007). Supervised machine learning: A review of classification techniques. *Emerging Artificial Intelligence Applications in Computer Engineering*, 3–24.
- [13]. Lakkaraju, H., Aguiar, E., Shan, C., Miller, D., Bhanpuri, N., Ghani, R., & Addison, K. L. (2015). A machine learning framework to identify students at risk of adverse academic outcomes. *Proceedings of the 21st ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, 1909–1918.
<https://doi.org/10.1145/2783258.2788620>
- [14]. McKinney, W. (2010). Data structures for statistical computing in Python. *Proceedings of the 9th Python in Science Conference*, 51–56.
- [15]. Mitchell, T. M. (1997). *Machine learning*. McGraw-Hill.
- [16]. Pedregosa, F., Varoquaux, G., Gramfort, A., et al. (2011). Scikit-learn: Machine learning in Python. *Journal of Machine Learning Research*, 12, 2825–2830.
- [17]. 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>
- [18]. Romero, C., Ventura, S., Pechenizkiy, M., & Baker, R. S. (2010). *Handbook of educational data mining*. CRC Press.
- [19]. Shmueli, G. (2010). To explain or to predict? *Statistical Science*, 25(3), 289–310.
<https://doi.org/10.1214/10-STS330>
- [20]. Siemens, G., & Baker, R. S. (2012). *Learning analytics and educational data mining: Towards communication and collaboration*. *Proceedings of the 2nd International Conference on Learning Analytics and Knowledge*, 252–254.
- [21]. Vapnik, V. N. (1998). *Statistical learning theory*. Wiley.
- [22]. Witten, I. H., Frank, E., Hall, M. A., & Pal, C. J. (2016). *Data mining: Practical machine learning tools and techniques* (4th ed.). Morgan Kaufmann.
- [23]. Zhou, Z.-H. (2012). *Ensemble methods: Foundations and algorithms*. CRC Press.