

**A NOVEL HYBRID FEATURE SELECTION FRAMEWORK FOR
ENHANCING ACCURACY AND INTERPRETABILITY IN MACHINE
LEARNING MODEL FOR STUDENT PERFORMANCE PREDICTION**

**Muhammad Sufyian Bin Mohd Azmi^{*1}, Mohd Hazli Bin Mohamed Zabil²,
Lim Kok Cheng³**

^{1,2,3}Department of Computing, College of Computing and Informatics, Universiti Tenaga
Selangor, Malaysia

Abstract

Accurate and interpretable student performance prediction machine learning model are crucial for identifying students at risk and improving academic outcomes in higher education. This study proposes a hybrid feature selection framework integrating Least Absolute Shrinkage and Selection Operator (LASSO) and Ant Colony Optimization (ACO) to enhance predictive accuracy and model interpretability at the same time. Utilizing the original dataset of 35 attributes, the proposed framework able to reduce dimensionality while retaining the most impactful features. LASSO identifies an initial subset of 12 features meanwhile ACO refines to an optimal six feature set in balancing accuracy and interpretability. The Random Forest model was trained on these features and achieved a remarkable accuracy of 99.21% and an AUC-ROC score of 0.9998, outperforming models using the full dataset, solely LASSO-selected features and other features selection techniques. This approach emphasizes the critical role of foundational knowledge, engagement and academic readiness in predicting student success. The proposed framework provides actionable insights for educators, enabling targeted interventions and fostering stakeholder trust in machine learning models for educational data mining

Key Words: Student performance prediction; Hybrid feature selection; Model interpretability; LASSO; Ant colony optimization; Educational data mining; Dimensionality reduction

1. Introduction

Student performance prediction is an important aspect of educational data mining (EDM), especially for faculty level courses like Programming 2 subject as it enables institutions to identify students at risk and provide targeted interventions to improve their academic outcomes [1]. Some students facing an academic challenges which can lead to high failure and withdrawal rates. Thus, this situation prompts a collaborative effort among teachers, parents and faculty management to address and overcome these academic challenges. Student academic performance serves as a key indicator that allows individuals to monitor their educational advancement, culminating in degree attainment or certification. Accurate prediction of student performance can help the management to identify students at risk [2], personalize their learning experiences and improve the academic outcomes. However, accurate prediction of student performance can be a significant challenge due to the multitude

of factors that influence it, such as demographic characteristics, prior academic achievement, socioeconomic status, and engagement with course materials [3]. On top of that, the biggest challenge lies in effectively selecting the most relevant features from large, complex datasets to build accurate and interpretable machine learning models.

Most researchers focus on feature selection to get high accuracy on machine learning models. Many of the existing research on student performance prediction in education focuses on optimizing accuracy through advanced feature selection techniques [4]. These techniques prioritize selecting the most predictive features, those that yield the highest accuracy scores for the model. Popular feature selection methods often emphasize accuracy, and include techniques such as Filter Methods, Wrapper Methods, and Embedded Methods [5]. Despite the success of these methods at improving predictive performance, focusing only on the accuracy often leads to a significant drawback which is a lack of interpretability in the final models [6]. While the model may have high accuracy, however the stakeholders such as lecturers, administrators and students themselves may not trust or fully understand the machine learning model make the decisions.

Studies showed that the behavior of many black box machine learning models whereby the decision-making process is hidden behind complex computations can make it difficult for these stakeholders to figure out on the reason of a particular decision was made [7]. This lack of transparency can lead to reduced confidence or trust in the model's predictions and making it harder to implement the model in real-world educational settings. Interpretability is really important in educational contexts because stakeholders need to understand the rationale behind the predictions to make informed on the decisions [8]. For example, if the machine learning model predicts that a student is at risk of failing the subject then educators need to know which factors contributed to this outcome to intervene effectively and at an early stage. Without an interpretability aspect, models may be dismissed as unreliable even if they perform well on standard accuracy metrics.

Our research proposes a novel hybrid feature selection framework aimed at striking a balance between accuracy and interpretability of machine learning models for predicting student performance in the Programming 2 course at Universiti Tenaga Nasional (UNITEN). By focusing not only on the model's predictive accuracy but also on how transparent and understandable its predictions are, we seek to build a framework that contributes to transparent models that stakeholders can trust and act upon. In doing so, we aim to contribute to a more explainable machine learning approach in educational data mining specifically for predicting student performance in the Programming 2 course at UNITEN.

2. Literature Review

The need for interpretable machine learning models in educational contexts has been highlighted in several studies. Study by [8] emphasizes that while machine learning models can achieve high predictive accuracy, the lack of interpretability hinders their acceptance and implementation in real-world educational settings.

2.1 Student Performance Prediction in Educational Data Mining

In recent years, EDM has gained significant attention to enhance decision-making in academic environments [9]. Predicting student performance is one of the priority tasks in EDM, where the goal is to predict students' future academic outcomes based on various aspect of the data such as past grades, attendance, assignments and engagement with the course content [10]. This prediction can help educators identify at-risk students early to tailor interventions and improve student learning outcomes. Machine learning models that are used for this task are leveraging large datasets that include academic records, engagement metrics, behavioral data and socio demographic features [11]. Many researchers have explored various machine learning techniques and approaches to predict student performance with a focus on getting high accuracy.

Thus, various machine learning techniques have been applied to predict student performance with a focus on getting high accuracy. According to [12], supervised learning algorithms such as Decision Trees (DT), Random Forests (RF), Support Vector Machines (SVM) and Neural Networks (NN) have been the most used. These methods offer strong predictive capabilities by learning patterns from historical datasets. However, the success of these models heavily depends on the quality and selection of features that best represent student behavior and performance [13].

Meanwhile, study by [14] showed DT are used due to their ability to provide both high accuracy on the prediction and fairly good interpretability aspect. The DT split the data into subsets based on the most significant features which makes them easy to understand and visualize. However, according to [15] DT can be prone to overfitting especially when using noisy or high-dimensional datasets. Random Forests (RF), an ensemble of DT have also shown promising results in student performance prediction as they can handle high-dimensional datasets and are less prone to overfitting [10]. However, like other ensemble methods, RF is considered as a black-box model, making it difficult to interpret [16]. Other techniques, such as Logistic Regression and Linear Regression, have also been used to predict student performance. While these models have demonstrated good predictive accuracy, but they also often lack interpretability [8] which is crucial in educational contexts where stakeholders need to understand the justification or reason behind the model predictions to make informed decisions. On top of that, even though these models are generally effective at prediction, however their ability to generalize well on unseen data and handle complex feature interactions depends heavily on how features are selected and processed. Thus, feature selection stage becomes an essential step to ensure that models can achieve both high accuracy and avoid overfitting [17].

2.2 Feature Selection Techniques in Machine Learning

Feature selection stage is one of the most important steps in the machine learning life cycle which can have a huge effect on model performance. In general, feature selection is one of the most important steps in the pre-processing of machine learning since it enhances model performance by removing irrelevant or redundant features from the dataset [18]. According to

[19], effective feature selection reduces model complexity, speeds up training and can enhance the generalization ability of the model. It is important in the context of student performance prediction where many features may be interrelated or noisy.

The most widely used feature selection techniques can be categorized into three groups, which are Filter Methods, Wrapper Methods and Embedded Methods [20]. Filter methods evaluate the relevance of each feature independently based on statistical measures, such as correlation, mutual information or chi-square test [21]. These methods are computationally efficient and can handle high-dimensional data, but on the other hand they do not consider the interaction between features [22]. In the context of student performance prediction, filter methods can help identify the most relevant factors, such as attendance, assignment grades, and demographic characteristics that contribute to student success. Wrapper methods use machine learning models as a black box to evaluate subsets of features and select the best performing subset [11]. These methods can capture feature interactions, but they are computationally more expensive and may be prone to overfitting [23]. Embedded methods integrate the feature selection process into the model training and allow the model to learn which features are most important [17]. Some common embedded methods include Lasso Regression (LR) in which it performs L1 regularization to encourage sparse feature selection and decision tree-based algorithms, which inherently perform feature selection during the tree-building process.

2.3 Challenge of Interpretability in Machine Learning Models

While improving accuracy through sophisticated feature selection methods is essential, it often comes at the expense of interpretability. According to [24] many machine learning models, such as Random Forests, Neural Networks and Gradient Boosting achieve state of the art performance but lack transparency. These complex machine learning models such as Random Forests, Neural Networks and Gradient Boosting, are often referred to as black box models, meaning it is difficult for stakeholders including educators, administrators and students to understand the specific reasons behind a model's predictions.

The inner workings and decision-making processes of these advanced models can be opaque, making it challenging for stakeholders to comprehend how the model arrived at a particular outcome.

2.3.1 The Importance of Interpretability

In the context of EDM, interpretability is important as educators, administrators and students need to understand the reasons behind the predictions to make informed decisions and interventions [8]. For example, if a model predicts that a student is likely to struggle in a course, stakeholders need to know which factors contributed most to that prediction in order to provide targeted support. Without interpretability, models are less likely to be trusted or adopted in real-world settings.

Current interpretability methods aim to provide insights into complex models are Shapley Values, LIME (Local Interpretable Model-agnostic Explanations) and Surrogate Models [24]. Shapley values originate from cooperative game theory [25]. Shapley values allocate each

features contribution to the prediction and offering a clear and interpretable explanation for individual predictions. Shapley values are particularly useful in ensemble methods like Random Forests, as they help explain the impact of each feature on a model's outcome [26]. LIME is another popular interpretability method that creates a local surrogate model to explain the predictions of a black-box model [27]. LIME is a technique that explains individual predictions by approximating the model's decision boundary locally with simpler and interpretable models. This helps users understand the decisions made by complex models on a case-by-case basis. Study by [24] showed surrogate models are simpler models for an examples decision trees or linear regression trained to approximate the predictions of more complex models. Surrogate models by approximating the behavior of the original black box model with simpler and more interpretable architectures also provide a means to understand the underlying reasoning behind the model's decisions.

Together these interpretability methods play a vital role in improving the transparency and reliability of complex machine learning models, especially in the context of educational data mining where stakeholder understanding is paramount.

2.4 Recent Advances in Feature Selection Techniques

Feature selection remains a dynamic area of research with several innovative methods emerging to address the limitations of traditional approaches. Table 1 below shows the comparison of recent feature selection techniques, focus and the achieved accuracy values on various datasets.

Table 1: Comparison of recent feature selection technique

Technique	Category	Focus	Interpretability Aspect	Typical Accuracy (%)
LASSO	Embedded method	Accuracy	Limited interpretability	85%–92%
Elastic Net	Embedded Method	Accuracy	Prioritizes multicollinearity handling over interpretability	88%–94%
Recursive Feature Elimination (RFE)	Wrapper Method	Accuracy	Model-dependent, making interpretability harder	90%–95%
Random Forest Feature Importance	Embedded Method	Accuracy	Often considered a black-box method, thus difficult to explain	92%–96%
Gradient Boosting (e.g., XGBoost)	Embedded Method	Accuracy	Provides importance rankings but lacks transparency	94%–97%

Chi-Square Test	Filter Method	Accuracy	Completely ignores feature interactions	75%–85%
Information Gain	Filter Method	Accuracy	Purely statistical; no explanation	80%–90%
Deep Learning Models (e.g., Neural Networks)	Feature Learning	Accuracy	Lacks interpretability due to the complex nature of deep networks	95%–99%
IGANN Sparse	Embedded Method	Balanced	Promotes sparsity through a non-linear feature selection process	88%–94%
UBayFS (User-Guided Bayesian Framework for Feature Selection)	Ensemble Method	Balanced	Integrates user guidance to enhance interpretability	87%–93%
Relief-Based Algorithms (e.g., ReliefF, SURF, MultiSURF)	Filter Method	Accuracy	While effective in identifying relevant features, they often lack direct interpretability	85%–92%
STIR (STatistical Inference Relief)	Filter Method	Accuracy	Enhances the statistical validity of feature selection but does not inherently improve interpretability.	88%–94%
Selective Bayesian Forest Classifier	Ensemble Method	Balanced	Provides a visualization tool that offers insight into relevant features	90%–96%

2.4.1 Detailed Analysis of Recent Advances Feature Selection Techniques

Most traditional and embedded methods such as Elastic Net, RFE, Random Forest, and Gradient Boosting prioritize accuracy. These methods are often evaluated purely on predictive performance, with little regard for how the selected features contribute to model transparency [28]. Accuracy for these methods typically falls in the 85%–99% range, depending on the dataset [29]. For example, Gradient Boosting achieves the highest accuracy

in the 94%–97% range [30], especially in complex datasets like credit scoring and recommendation systems. Elastic Net is effective in high-dimensional and sparse datasets achieving 85%–94% accuracy but offer limited transparency [31].

Recent frameworks like IGANN Sparse, UBayFS, and Selective Bayesian Forest Classifier attempt to combine accuracy with interpretability by promoting sparsity, integrating user guidance, or providing visualization tools [32]. These methods deliver balanced accuracy in the range of 87%–96% while offering insights into feature importance and interactions. Techniques like Chi-Square Test and Information Gain are fast and simple and suitable for preprocessing large datasets [33]. However, their accuracy is lower in between 75%–90% and they lack mechanisms to detect feature interactions or improve interpretability. These methods are best used as a first step followed by more advanced techniques to refine feature selection. RFE, Random Forest Feature Importance and Gradient Boosting dominate practical applications due to their high accuracy in the range of 90%–97%. These methods are particularly effective for tasks requiring complex feature interactions. However, these techniques are criticized for their black-box nature and making it difficult to justify predictions to non-technical stakeholders [34].

Techniques like LASSO and IGANN Sparse achieve both accuracy and interpretability by selecting a minimal set of features. Sparse modelling is gaining traction for applications in genomics, drug discovery and personalized medicine where too many features can overwhelm decision making. Accuracy for sparse techniques ranges from 85%–94%, slightly lower than ensemble methods but more interpretable. Frameworks like UBayFS integrate user expertise into feature selection, enhancing interpretability by aligning feature relevance with domain knowledge. These methods are particularly useful in personalized medicine and marketing where stakeholders need actionable insights. Techniques like the Selective Bayesian Forest Classifier use visualization tools to represent feature importance and interactions, making them highly interpretable. These methods maintain accuracy in the range of 90%–96% while addressing stakeholder concerns about transparency.

2.5 Gap in Current Literature

The literature on student performance prediction largely focuses on improving accuracy through advanced feature selection and model tuning. However, there is a notable lack of studies that integrate interpretability alongside accuracy [35]. In reality, machine learning models that are optimized solely for accuracy may achieve high performance but fail to offer meaningful explanations for their predictions. This makes it difficult for stakeholders to trust and act upon the results, especially in educational contexts [36].

For example, while Random Forests and Gradient Boosting are powerful tools for prediction, they are often criticized for being unclear [24]. According to [37], neural network models particularly deep learning models, offer even greater predictive power but are typically seen as black-box models. These models are difficult to interpret and their use in real world educational environments is limited because stakeholders need to understand why a model made a specific prediction such as why a student is predicted to fail.

This gap presents an opportunity to develop a new feature selection framework that not only optimizes accuracy but also enhances the interpretability of the model. Such a framework would address the concerns of stakeholders by providing transparent and understandable explanations for why specific predictions are made.

3. Methodology

To address the limitations in the current literature, this section proposes a new hybrid feature selection framework that aims to increase both the accuracy and interpretability of machine learning models for student performance prediction. In order to conduct the experimental study, the real datasets were collected from the UNITEN student performance database for the Programming 2 subject. For this study, the data encompassing various dimensions of student records. The experiment was conducted using Jupyter Notebook, a popular interactive computing environment. Jupyter Notebook allows for the integration of code, visualizations, and explanatory text, facilitating a collaborative and iterative approach to data analysis and model development [38].

3.1. Datasets

The dataset used in this study contained a comprehensive set of 35 attributes, capturing demographic, academic, behavioral and socio-economic details of students as shown in Table 2. These attributes were chosen for their potential relevance in predicting student performance in the Programming 2 course. These features were initially considered relevant for the prediction task due to their potential influence on student performance in the Programming 2 course.

Table 2: Uniten dataset attributes

Category	Attributes
Demographic Data	Date Of Birth, Gender, Nationality, State, Country Name
Socio-Economic Data	Monthly Income, Income Range, Sponsor, Sponsor Category, Amount, Sponsor Type
Academic Background	Entry Qualification, Level, Program Name, Semester Attended, Tahun Pengajian
Attendance	Problem Solving Attendance Percentage, Programming 1 Attendance Percentage,
	Computer Organization Attendance Percentage, Discrete Structure Attendance Percentage,
	Programming 2 Attendance Percentage,

	Attendance Percentage
Grades	Programming 1 Grade, Computer Organization Grade, Discrete Structure Grade, Problem Solving Grade
Other Attributes	Mode Of Study, College, Campus, Intake, Credit Unit, Section, Date Registered

3.2 Handling Missing Values

Missing values were identified in the UNITEN dataset and addressed systematically to ensure data completeness and avoid biases introduced by incomplete records. The following steps were undertaken:

- Missing values were identified in numerical and categorical columns. Features with a significant percentage of missing values (>30%) were reviewed for potential exclusion but no features were removed as all were deemed essential.
- Missing values in numerical features such as subject scores and attendance percentage were replaced with the mean value of the respective feature. This ensures that the overall distribution of the feature remains consistent while missing values in categorical features such as gender and entry qualification were imputed using the mode of the respective feature. This approach ensures minimal disruption to the distribution of the categorical data [39].

The mean and mode imputation techniques were chosen for their simplicity and effectiveness, particularly in datasets with relatively small missing rates [40]. This allowed us to preserve the underlying data structure while improving the completeness of the dataset.

3.3 One-Hot Encoding for Categorical Features

To enable compatibility with machine learning algorithms like LASSO and to retain the information embedded in categorical features, all categorical variables were transformed into numerical representations using one-hot encoding. This method creates binary indicator columns for each category in ensuring no ordinal relationships are assumed in the data. By representing categorical variables in a numerical format, the feature selection algorithms will effectively evaluate their importance in the predictive modelling process [41]. One-hot encoding was chosen as it preserves categorical information without introducing unintended ordinal relationships [42]. Thus, it will ensure that each category is treated independently, which is essential for linear models like LASSO.

3.4 Proposed Hybrid Feature Selection Framework

The proposed hybrid framework is integrating LASSO (Least Absolute Shrinkage and Selection Operator) and Ant Colony Optimization (ACO) to create a robust feature selection process that balances predictive accuracy and interpretability. By leveraging the strengths of both techniques, the framework systematically identifies and optimizes a subset of relevant

features to ensure high performance while maintaining simplicity.

Feature selection is critical in predictive modelling as it reduces dimensionality so it will enhance model performance and increase interpretability. Existing methods often prioritize high accuracy over interpretability and making it difficult for stakeholders to understand the decision-making process of machine learning models. The proposed framework addresses these limitations by using LASSO for initial feature reduction. By doing that, it is ensuring sparsity and eliminating irrelevant features. On top of that, employing ACO to optimize the reduced feature set based on a fitness function that balances accuracy and interpretability.

This hybrid approach is particularly suitable for educational datasets, where transparency is essential for gaining the trust of educators and stakeholders. The workflow of the proposed framework consists of three main stages which are Initial Feature Reduction, Feature Optimization and Model Training and Evaluation. Illustration of the structure of the proposed feature selection framework is shown in Figure 1 below.

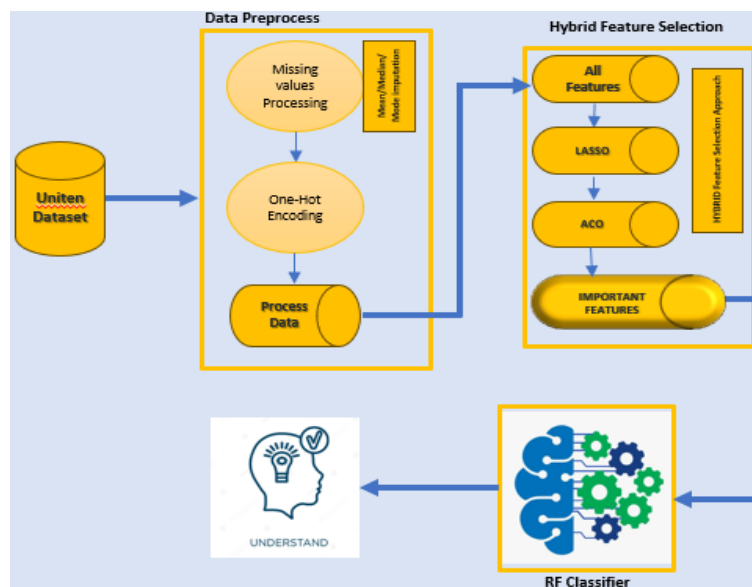


Figure 1: Structure of the proposed feature selection framework

3.4.1 Stage 1: Initial Feature Reduction with Lasso

LASSO performs feature selection by minimizing a loss function that incorporates an L1 regularization term. This term penalizes the absolute values of the feature coefficients, effectively shrinking irrelevant feature coefficients to zero, resulting in a sparse model. The LASSO objective function is given by in Equation (1).

$$\hat{\beta} = \underset{\beta}{argmin} \left[\frac{1}{2n} \sum_{i=1}^n \left(y_i - \sum_{j=1}^p X_{ij} \beta_j \right)^2 + \lambda \sum_{j=1}^p |B_j| \right]$$

The initial dataset consisted of a feature matrix that included all 35 attributes as detailed earlier. These attributes spanned various categories including demographic data, socio-

economic indicators, academic background, attendance records and grades. To perform feature selection, the LASSO utilizes a regularization parameter $\lambda=0.01$ to control sparsity. This parameter penalized irrelevant or redundant feature coefficients, effectively shrinking them to zero, while retaining only the most predictive attributes. Through this process, features with non-zero coefficients were identified as relevant for predicting the target variable, which was the 1 or 0 status in Programming 2. Conversely, features with coefficients equal to zero were excluded, as they contributed insignificantly to the model. The resulting subset of attributes was then used for further optimization and modelling. Fig. 2 below showed that each curve represents a coefficient, and the x-axis represents the regularization penalty parameter. As λ changes, a coefficient that becomes non-zero enters the LASSO regression model [43].

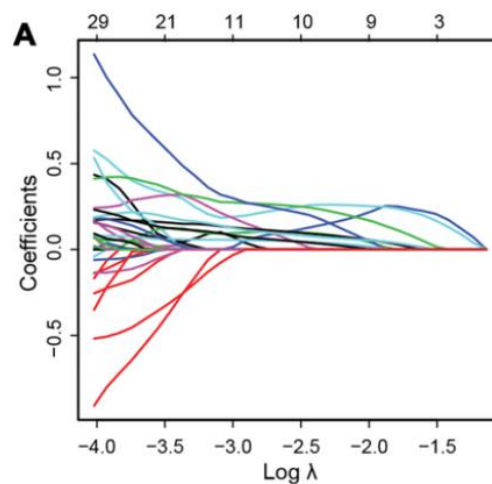


Figure 2: LASSO regression model

The LASSO feature selection process reduced the dataset from 35 attributes to 12 key attributes and significantly improving interpretability and computational efficiency. The retained attributes are shown in Table 3. This subset of features retains variables most predictive of student performance while excluding those that contribute little to the model.

Table 3: Attributes retained after LASSO

Category	Attributes
Demographic Data	EntryQualification
Socio-Economic Data	Income Range, Sponsor Type
Academic Background	SemesterAttended
Attendance	Problem Solving Attendance Percentage, Programming 1 Attendance Percentage,
	Computer Organization Attendance Percentage, Discrete Structure Attendance Percentage,

	Programming 2 Attendance Percentage
Grades	Programming 1 Grade, Computer Organization Grade, Discrete Structure Grade

3.4.2 Stage 2: Feature Optimization with ACO

After the initial feature reduction using LASSO, the retained features were further optimized using ACO to identify the best subset of features. ACO is a nature-inspired metaheuristic algorithm that simulates the foraging behaviour of ants to find optimal solutions [44]. In this study, ACO was employed to explore and evaluate different combinations of the LASSO-selected features, balancing predictive accuracy and interpretability through a fitness function. The optimization process involved the following steps:

- A colony of ants was initialized, each representing a possible subset of features.
- Pheromone trails were assigned equally across all LASSO-selected features at the start of the process.
- Each ant traversed a feature graph, where nodes represented features, and edges represented the probability of selecting a specific feature based on pheromone intensity and heuristic information.
- The probability of an ant selecting feature at time calculated as in Equation (2) below.

$$P_{ij} = \frac{\tau_{ij}^{\alpha} \cdot n_{ij}^{\beta}}{\sum_{k \in \text{Features}} \tau_{ik}^{\alpha} \cdot n_{ik}^{\beta}} \quad (2)$$

where P_{ij} is probability of selecting feature j , τ_{ij} is a pheromone level for feature j , n_{ij} is a feature importance score from LASSO and α, β is a control parameter that determine the influence of pheromone intensity and heuristic information.

- Then the next step is fitness evaluation where each ant constructed a feature subset and evaluated its performance using a fitness function as in Equation (3) below.

$$\text{Fitness} = \alpha(\text{Accuracy}) + \beta \left(\frac{1}{\text{Number of features}} \right) \quad (3)$$

where α, β is a weight that balance accuracy and interpretability in the optimization process.

- The following step is pheromone update where pheromone levels were updated based on the performance of feature subsets constructed by the ants. The pheromone update rule was as in Equation (4) below.

$$\tau_{ij} = (1 - \rho) \cdot \tau_{ij} + \Delta\tau_{ij} \quad (4)$$

where ρ is a pheromone evaporation rate, which prevents premature convergence and $\Delta\tau_{ij}$ is

contribution of the best performing subsets to the pheromone level.

- The process was repeated until the algorithm converged to an optimal feature subset.
- Lastly, the optimal subset of features was selected based on the highest fitness value.

After applying ACO to the feature set refined by LASSO, the algorithm identified the optimal subset of features for predicting student performance in the Programming 2 course. These features were selected based on their ability to maximize predictive accuracy while minimizing redundancy in the dataset. The best feature subset selected by ACO are shown in Table 4.

Table 4: Attributes retained after LASSO+ACO

Category	Attributes
Demographic Data	Entry Qualification
Academic Background	Semester Attended
Attendance	Computer Organization Attendance Percentage,
	Programming 2 Attendance Percentage
Grades	Programming 1 Grade, Computer Organization Grade

These six features were identified as the most relevant predictors of student performance based on the ACO optimization process. Together, they balance on academic performance, behavioral engagement and background information.

3.5 Model Training and Evaluation

The final subset of features selected by ACO was used to train and evaluate a predictive model for student performance in the Programming 2 course. A Random Forest (RF) classifier was chosen for this task due to its robustness and ability to handle feature interactions effectively [45]. The model's performance was assessed using metrics such as accuracy, precision, recall, F1-score, and AUC-ROC.

RF classifier was selected for evaluating and validating the final subset of features identified by ACO due to its robustness, versatility, and suitability for this study. As an ensemble method, Random Forest combines predictions from multiple decision trees, making it highly resistant to overfitting and capable of handling noisy data [16]. This robustness ensures consistent performance across varied feature types, such as numerical grades, attendance

percentages, and categorical variables like entry qualifications.

4. Results

The results of this study demonstrate the effectiveness of the proposed hybrid feature selection framework, which hybrid of LASSO and ACO to enhance both predictive accuracy and interpretability for student performance prediction in the Programming 2 course. The outcomes are discussed below in terms of feature selection, model performance, and implications.

4.1 Feature Selection Results

The initial dataset consisted of 35 attributes, representing a comprehensive set of demographics, academic, and behavioural variables. After applying LASSO for initial feature reduction, the number of attributes was reduced to 12, retaining only those with non-zero coefficients. Subsequently, ACO further refined the feature subset by exploring various combinations of the LASSO selected features and evaluating their predictive performance as shown in Figure 2. The final subset of 6 features selected by ACO as shown in Table 3. This optimized subset balances high predictive power with interpretability and focusing on academic performance, engagement and preparedness.

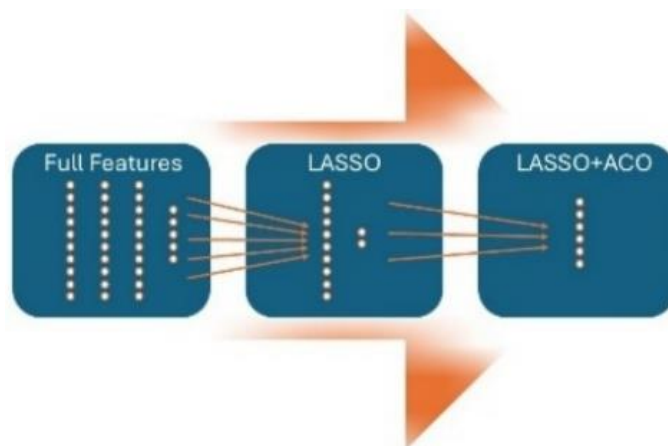


Figure 3: Features reduction after proposed hybrid framework

To validate the efficiency of the proposed framework, experiments were conducted with other widely use feature selection techniques, such as Recursive Feature Elimination (RFE), Principal Component Analysis (PCA), Correlation Based Feature Selection (CFS), Chi Square, Gain Ratio (GR), Information Gain (IG) and Random Forest-based feature importance. The result of feature selection reduction is shown in Table 5 below.

Table 5: Number of features selected by each feature selection algorithm

Feature Selection Technique	Number of features
RFE	18

PCA	10
CFS	10
Chi Square	15
GR	15
IG	15
Random Forest Based	20
LASSO	12
Proposed Framework	6

While these methods achieved varying degrees of feature reduction, none provided the same level of balance between dimensionality reduction and predictive accuracy as the LASSO + ACO framework. Based on the performance analysis of different feature selectors from Tables 5, the following observations are made:

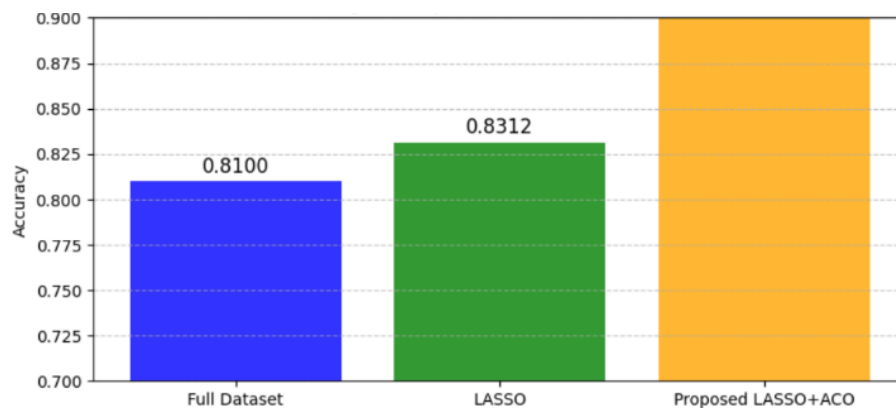
- LASSO + ACO framework reduced the number of features to just 6 representing a 82.86% reduction.
- RFE and Random Forest retained 18 and 20 features respectively, resulting in higher model complexity.

LASSO + ACO framework reduced the number of features to 82.86% reduction, the highest among all tested methods, while preserving essential predictors directly tied to academic performance and engagement.

4.2. Model Performance with Proposed Framework Hybrid Features

The performance of the proposed hybrid feature selection framework, integrating LASSO and ACO, was evaluated by training a Random Forest Classifier on the final subset of features identified through this process. This hybrid approach effectively reduced the original dataset to just six key attributes, preserving only the most critical predictors of student performance. The evaluation focused on metrics such as accuracy, AUC-ROC, precision, recall, and F1-score to assess the model's ability to generalize and make reliable predictions.

To validate the effectiveness of our proposed framework, a comparative analysis was conducted to evaluate the model's accuracy across different feature subsets which are the full dataset, the LASSO-reduced feature set, and the LASSO + ACO proposed hybrid feature selection framework in Figure 3. While the details evaluation is presented in Table 6 – Table 8 below.

**Figure 4:** RF classifier accuracy comparison across feature subsets**Table 6:** Model performance metrics using all attributes

Metric	Class 0 (Fail)	Class 1 (Pass)	Macro Average	Weighted Average	Overall
Precision	0.93	0.77	0.85	0.83	-
Recall	0.5	0.98	0.74	0.81	-
F1-Score	0.65	0.86	0.76	0.79	-
Support	28	49	-	-	77
Accuracy	-	-	-	-	0.81
AUC-ROC	-	-	-	-	0.9125

Table 7: Model performance metrics using LASSO selected attributes

Metric	Class 0 (Fail)	Class 1 (Pass)	Macro Average	Weighted Average	Overall
Precision	1	0.79	0.9	0.87	-
Recall	0.54	1	0.77	0.83	-
F1-Score	0.7	0.88	0.79	0.82	-
Support	28	49	-	-	77
Accuracy	-	-	-	-	0.83
AUC-ROC	-	-	-	-	0.9009

Table 8: Model performance metrics using proposed framework selected attributes

Metric	Class 0 (Fail)	Class 1 (Pass)	Macro Average	Weighted Average	Overall
Precision	0.99	0.99	0.99	0.99	-
Recall	0.98	1	0.99	0.99	-
F1-Score	0.99	0.99	0.99	0.99	-
Support	106	275	-	-	381
Accuracy	-	-	-	-	0.99
AUC-ROC	-	-	-	-	0.9998

The results in Table 8 above indicate exceptional performance with the model correctly predicting 99.21% of student outcomes. The near-perfect AUC-ROC with score 0.9998 demonstrates the model's excellent ability to distinguish between students who pass and fail.

In Table 9 below, we see further evidence that the proposed hybrid feature selection framework outperforms other feature selection techniques on RF classifier in terms of accuracy of RF classifier.

Table 9: Accuracy of RF on selected features for each feature selection algorithm

Feature Selection Technique	RF Model Accuracy
RFE	0.85
PCA	0.80
CFS	0.75
Chi Square	0.86
GR	0.70
IG	0.72
Random Forest Based	0.78
LASSO	0.83
Proposed Framework	0.99

5. Discussion

We proposed the new hybrid LASSO and ACO feature selection framework because this hybrid approach leverages the strengths of both techniques. LASSO efficiently reduces high-dimensional datasets to a sparse subset of predictive features while maintaining interpretability [46-47]. ACO further optimizes this reduced subset by exploring feature

combinations and addressing non-linear relationships [48]. Other combinations, such as RFE + ACO or PCA + ACO, either introduce additional computational complexity or sacrifice interpretability, making them less suited to the goals of this study. The LASSO + ACO framework achieves an ideal balance between dimensionality reduction, interpretability, and predictive accuracy, as demonstrated by our experimental results.

The results validate the effectiveness of the LASSO + ACO hybrid framework in identifying an optimal subset of features. By integrating LASSO for dimensionality reduction and sparsity, the framework ensures that irrelevant or redundant features are eliminated, simplifying the dataset. ACO complements this process by fine-tuning the feature selection, enabling the identification of a subset that enhances predictive accuracy while maintaining interpretability.

The high accuracy and AUC-ROC scores achieved by the model suggest that the selected features effectively capture critical determinants of student performance. These include prior academic performance, attendance and readiness which are essential indicators for predicting success in the Programming 2 course.

By reducing the feature set to six attributes, the framework significantly improves interpretability. This reduction allows educators and stakeholders to clearly understand the key factors influencing student outcomes. For instance, the inclusion of Programming 1 Grade underscores the importance of foundational knowledge, while the focus on Attendance Percentages highlights the role of consistent engagement in the classroom.

Furthermore, the results underscore the potential of data driven approaches for early intervention. The identified features provide actionable insights, enabling educators to design targeted support strategies. For example, improving attendance rates or offering additional resources to students with lower grades can address critical gaps and enhance student performance outcomes.

6. Conclusion And Future Work

This study proposed and validated a hybrid feature selection framework integrating LASSO and Ant Colony Optimization (ACO) for predicting student performance in the Programming 2 course at UNITEN. The hybrid framework effectively reduced the original dataset of 35 attributes to a concise subset of 6 key features, striking a balance between predictive accuracy and interpretability. The integration of LASSO ensured dimensionality reduction by removing irrelevant features, while ACO refined the subset by exploring feature combinations to maximize accuracy and simplicity.

The results demonstrated the effectiveness of the proposed framework, with the Random Forest model trained on the ACO-selected features achieving an impressive accuracy of 99.21% and a near-perfect AUC-ROC score of 0.9998. These findings highlight that the selected features, which included academic performance, attendance records, and entry qualifications, captured the critical determinants of student success. The framework not only enhanced predictive performance but also provided actionable insights for educators and stakeholders, making it a valuable tool for designing early interventions to improve student

outcomes.

By improving interpretability through a reduced feature set, the framework addressed the common challenge of stakeholder trust in machine learning models. Educators can now better understand and leverage key predictive factors, such as foundational grades and engagement levels, to identify at-risk students and implement targeted support strategies.

Future work could explore the application of this framework to larger datasets or other courses to evaluate its generalizability. Additionally, integrating advanced interpretability tools, such as SHAP [49] values or LIME [50] could provide further transparency regarding the model's decision-making process. Overall, the proposed LASSO + ACO hybrid framework presents a robust, interpretable, and effective approach to student performance prediction, contributing significantly to the field of educational data mining.

7. Acknowledgment

We would like to express our sincere gratitude to YCU (Yayasan Canselor Universiti Tenaga Nasional) and UNITEN (Universiti Tenaga Nasional) for their generous support through the YCU COMMUNITY ENERGY GRANT 2025 under project code YCUCE2025006. This sponsorship has played a pivotal role in facilitating the research presented in this paper, and we deeply appreciate their commitment to advancing academic and scientific endeavours.

References

- [1] A. Tamhane, S. Ikbal, B. Sengupta, M. Duggirala, and J. Appleton, "Predicting student risks through longitudinal analysis," in *Proceedings of the 20th ACM SIGKDD international conference on Knowledge discovery and data mining*, New York, NY, USA: ACM, Aug. 2014, pp. 1544–1552. doi: 10.1145/2623330.2623355.
- [2] C. Lei and K. F. Li, "Academic Performance Predictors," in *2015 IEEE 29th International Conference on Advanced Information Networking and Applications Workshops*, IEEE, Mar. 2015, pp. 577–581. doi: 10.1109/WAINA.2015.114.
- [3] A. Namoun and A. Alshanqiti, "Predicting Student Performance Using Data Mining and Learning Analytics Techniques: A Systematic Literature Review," *Applied Sciences*, vol. 11, no. 1, p. 237, Dec. 2020, doi: 10.3390/app11010237.
- [4] X. Ma and Z. Zhou, "Student pass rates prediction using optimized support vector machine and decision tree," in *2018 IEEE 8th Annual Computing and Communication Workshop and Conference (CCWC)*, IEEE, Jan. 2018, pp. 209–215. doi: 10.1109/CCWC.2018.8301756.
- [5] B. Butcher and B. J. Smith, "Feature Engineering and Selection: A Practical Approach for Predictive Models," *Am Stat*, vol. 74, no. 3, pp. 308–309, Jul. 2020, doi: 10.1080/00031305.2020.1790217.
- [6] W. J. Murdoch, C. Singh, K. Kumbier, R. Abbasi-Asl, and B. Yu, "Definitions, methods, and applications in interpretable machine learning," *Proceedings of the*

- National Academy of Sciences*, vol. 116, no. 44, pp. 22071–22080, Oct. 2019, doi: 10.1073/pnas.1900654116.
- [7] D. V. Carvalho, E. M. Pereira, and J. S. Cardoso, “Machine Learning Interpretability: A Survey on Methods and Metrics,” *Electronics (Basel)*, vol. 8, no. 8, p. 832, Jul. 2019, doi: 10.3390/electronics8080832.
- [8] M. Chitti, P. Chitti, and M. Jayabalan, “Need for Interpretable Student Performance Prediction,” in *2020 13th International Conference on Developments in eSystems Engineering (DeSE)*, IEEE, Dec. 2020, pp. 269–272. doi: 10.1109/DeSE51703.2020.9450735.
- [9] J. Cheng, “Data-Mining Research in Education,” Mar. 2017, [Online]. Available: <http://arxiv.org/abs/1703.10117>
- [10] H. K. R. Toppireddy, B. Saini, and P. S. Hada, “Academic Enhancement System using EDM Approach,” in *2019 Second International Conference on Advanced Computational and Communication Paradigms (ICACCP)*, IEEE, Feb. 2019, pp. 1–10. doi: 10.1109/ICACCP.2019.8882954.
- [11] H. Lakkaraju *et al.*, “A Machine Learning Framework to Identify Students at Risk of Adverse Academic Outcomes,” in *Proceedings of the 21th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, New York, NY, USA: ACM, Aug. 2015, pp. 1909–1918. doi: 10.1145/2783258.2788620.
- [12] V. Vijayalakshmi and K. Venkatachalapathy, “Comparison of Predicting Student’s Performance using Machine Learning Algorithms,” *International Journal of Intelligent Systems and Applications*, vol. 11, no. 12, pp. 34–45, Dec. 2019, doi: 10.5815/ijisa.2019.12.04.
- [13] J. L. Rastrollo-Guerrero, J. A. Gómez-Pulido, and A. Durán-Domínguez, “Analyzing and Predicting Students’ Performance by Means of Machine Learning: A Review,” *Applied Sciences*, vol. 10, no. 3, p. 1042, Feb. 2020, doi: 10.3390/app10031042.
- [14] S. Herzog, “Estimating student retention and degree-completion time: Decision trees and neural networks vis-à-vis regression,” *New Directions for Institutional Research*, vol. 2006, no. 131, pp. 17–33, Sep. 2006, doi: 10.1002/ir.185.
- [15] S. Jayaprakash, S. Krishnan, and V. Jaiganesh, “Predicting Students Academic Performance using an Improved Random Forest Classifier,” in *2020 International Conference on Emerging Smart Computing and Informatics (ESCI)*, IEEE, Mar. 2020, pp. 238–243. doi: 10.1109/ESCI48226.2020.9167547.
- [16] V. Svetnik, A. Liaw, C. Tong, J. C. Culberson, R. P. Sheridan, and B. P. Feuston, “Random Forest: A Classification and Regression Tool for Compound Classification and QSAR Modeling,” *J Chem Inf Comput Sci*, vol. 43, no. 6, pp. 1947–1958, Nov. 2003, doi: 10.1021/ci034160g.
- [17] J. Li and H. Liu, “Challenges of Feature Selection for Big Data Analytics,” Nov. 2016,

[Online]. Available: <http://arxiv.org/abs/1611.01875>

- [18] T. Amr, B. De, and L. Iglesia, “Survey on Feature Selection,” 2015.
- [19] B. Ghogh et al., “Feature Selection and Feature Extraction in Pattern Analysis: A Literature Review,” May 2019, [Online]. Available: <http://arxiv.org/abs/1905.02845>
- [20] G. Chandrashekar and F. Sahin, “A survey on feature selection methods,” *Computers & Electrical Engineering*, vol. 40, no. 1, pp. 16–28, Jan. 2014, doi: 10.1016/j.compeleceng.2013.11.024.
- [21] A. Kaur, K. Guleria, and N. Kumar Trivedi, “Feature Selection in Machine Learning: Methods and Comparison,” in *2021 International Conference on Advance Computing and Innovative Technologies in Engineering (ICACITE)*, IEEE, Mar. 2021, pp. 789–795. doi: 10.1109/ICACITE51222.2021.9404623.
- [22] G. Sosa-Cabrera, S. Gómez-Guerrero, M. García-Torres, and C. E. Schaerer, “Feature Selection: A perspective on inter-attribute cooperation,” Jun. 2023, doi: 10.1007/s41060-023-00439-z.
- [23] R. Kohavi and G. H. John, “Wrappers for feature subset selection,” *Artif Intell*, vol. 97, no. 1–2, pp. 273–324, Dec. 1997, doi: 10.1016/S0004-3702(97)00043-X.
- [24] M. T. Ribeiro, S. Singh, and C. Guestrin, ““Why should i trust you?” Explaining the predictions of any classifier,” in *Proceedings of the ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, Association for Computing Machinery, Aug. 2016, pp. 1135–1144. doi: 10.1145/2939672.2939778.
- [25] L. Merrick and A. Taly, “The Explanation Game: Explaining Machine Learning Models Using Shapley Values,” Sep. 2019, [Online]. Available: <http://arxiv.org/abs/1909.08128>
- [26] D. Basu, “On Shapley Credit Allocation for Interpretability,” Dec. 2020, [Online]. Available: <http://arxiv.org/abs/2012.05506>
- [27] K. Aas, M. Jullum, and A. Løland, “Explaining individual predictions when features are dependent: More accurate approximations to Shapley values,” *Artif Intell*, vol. 298, p. 103502, Sep. 2021, doi: 10.1016/j.artint.2021.103502.
- [28] C. Luan and G. Dong, “Experimental Identification of Hard Data Sets for Classification and Feature Selection Methods with Insights on Method Selection.”
- [29] J. L. Speiser, M. E. Miller, J. Tooze, and E. Ip, “A comparison of random forest variable selection methods for classification prediction modeling,” *Expert Syst Appl*, vol. 134, pp. 93–101, Nov. 2019, doi: 10.1016/j.eswa.2019.05.028.
- [30] C. Bentéjac, A. Csörgő, and G. Martínez-Muñoz, “A Comparative Analysis of XGBoost,” Nov. 2019, doi: 10.1007/s10462-020-09896-5.
- [31] R. Odegua, “Predicting Bank Loan Default with Extreme Gradient Boosting.”
- [32] H. Lakkaraju, S. H. Bach, and J. Leskovec, “Interpretable Decision Sets,” in

- Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, New York, NY, USA: ACM, Aug. 2016, pp. 1675–1684. doi: 10.1145/2939672.2939874.
- [33] A. Karper, “Feature and Variable Selection in Classification,” Feb. 2014, [Online]. Available: <http://arxiv.org/abs/1402.2300>
- [34] O. Loyola-Gonzalez, “Black-Box vs. White-Box: Understanding Their Advantages and Weaknesses From a Practical Point of View,” *IEEE Access*, vol. 7, pp. 154096–154113, 2019, doi: 10.1109/ACCESS.2019.2949286.
- [35] R. Alamri and B. Alharbi, “Explainable Student Performance Prediction Models: A Systematic Review,” *IEEE Access*, vol. 9, pp. 33132–33143, 2021, doi: 10.1109/ACCESS.2021.3061368.
- [36] D. Petkovic, “It is not ‘accuracy vs. explainability’-we need both for trustworthy AI systems.”
- [37] C. Rudin, “Stop Explaining Black Box Machine Learning Models for High Stakes Decisions and Use Interpretable Models Instead,” Nov. 2018, [Online]. Available: <http://arxiv.org/abs/1811.10154>
- [38] J. M. Perkel, “Why Jupyter is data scientists’ computational notebook of choice,” *Nature*, vol. 563, no. 7729, pp. 145–146, Nov. 2018, doi: 10.1038/d41586-018-07196-1.
- [39] P. Mishra, K. Divya Mani, P. Johri, and D. Arya, “FCMI: Feature Correlation based Missing Data Imputation.”
- [40] A. Farhangfar, L. Kurgan, and J. Dy, “Impact of imputation of missing values on classification error for discrete data,” *Pattern Recognit*, vol. 41, no. 12, pp. 3692–3705, Dec. 2008, doi: 10.1016/j.patcog.2008.05.019.
- [41] Max Kuhn and Kjell Johnson, *Feature Engineering and Selection: A Practical Approach for Predictive Models*. CRC Press, 2019.
- [42] K. Potdar, T. S., and C. D., “A Comparative Study of Categorical Variable Encoding Techniques for Neural Network Classifiers,” *Int J Comput Appl*, vol. 175, no. 4, pp. 7–9, Oct. 2017, doi: 10.5120/ijca2017915495.
- [43] J. Yan *et al.*, “Development of a four-gene prognostic model for pancreatic cancer based on transcriptome dysregulation,” *Aging*, vol. 12, no. 4, pp. 3747–3770, Feb. 2020, doi: 10.18632/aging.102844.
- [44] S. Kashef and H. Nezamabadi-pour, “An advanced ACO algorithm for feature subset selection,” *Neurocomputing*, vol. 147, pp. 271–279, Jan. 2015, doi: 10.1016/j.neucom.2014.06.067.
- [45] Ryan A. Estrellado, Emily A. Freer, Joshua M. Rosenberg, and sabella C. Velásquez, *Data Science in Education Using R*, 2nd Edition. 2020.

- [46] Shailabh Verma, “Unlocking the Power of LASSO Regression,” pickl.ai.
- [47] Z. Wang, X. Xiao, and S. Rajasekaran, “Novel and efficient randomized algorithms for feature selection,” *Big Data Mining and Analytics*, vol. 3, no. 3, pp. 208–224, Sep. 2020, doi: 10.26599/BDMA.2020.9020005.
- [48] M. Deriche, “Feature Selection using Ant Colony Optimization,” in *2009 6th International Multi-Conference on Systems, Signals and Devices*, IEEE, Mar. 2009, pp. 1–4. doi: 10.1109/SSD.2009.4956825.
- [49] V. Swamy, B. Radmehr, N. Krco, M. Marras, and T. Käser, “Evaluating the Explainers: Black-Box Explainable Machine Learning for Student Success Prediction in MOOCs,” Jul. 2022, [Online]. Available: <http://arxiv.org/abs/2207.00551>
- [50] V. Swamy, S. Du, M. Marras, and T. Kaser, “Trusting the Explainers: Teacher Validation of Explainable Artificial Intelligence for Course Design,” in *ACM International Conference Proceeding Series*, Association for Computing Machinery, Mar. 2023, pp. 345–356. doi: 10.1145/3576050.3576147.