

Machine Learning-Based Predictive Assessment in Teaching and Learning

Dr. Pravin Ramesh Gundalwar,

Professor, SOCSE, Sandip University Nashik

Abstract— The main objective of engineering education has expanded enormously to provide a quality education to its students. Determining the knowledge for prediction about student performance in engineering education is one of the achievements of any such institute. Student performance can be analyzed and predicted by using machine learning algorithms on some available data sets. This paper focuses on the capabilities of selected machine learning algorithms in the context of engineering education by offering a proposed model for the engineering education institute system. Naïve Bayes, Random Forest, Decision Stump, K Star and Decision Tree algorithms are used in assessment analysis in the teaching and learning process. These algorithms predict the student's expected performance based on one of the identified attributes. The percentage split data validation technique is used for training and testing on the WEKA platform. The analysis based on performance metrics is useful for predicting and improving the student's next assessment based on their existing performance.

Keywords—Machine learning algorithm, engineering education, WEKA, Naïve Bayes, Decision Tree

I. INTRODUCTION

This study aims to contribute to engineering education institutions and students for their performance in any online/offline assessment test. This would help better educationally by analyzing student scores and identifying those needing help with their weak section in offered studies. It would be ideal if a machine learning prediction-based modeled platform could be created that can help all the stakeholders of engineering education institutions viz., staff, students, parents, and recruiters. Rapid assessment tests and their analysis are a need of all education institutions around the world which can be an effective tool for assessment, feedback, and improvement areas.

Machine Learning (ML) can be an irresistible domain to tackle as it encompasses a large subset of research and studies on different algorithms to solve complex tasks. The process of learning about a student's knowledge and comprehension of a particular course is referred to as student knowledge assessment. ML algorithms are used to predict students' performance and assess their knowledge by conducting assessment tests (Anbukarasi V. et al., 2019)

The main functions of ML and its application methods using algorithms are to discover and extract patterns of stored data. Data mining and knowledge discovery applications have got a rich focus due to their significance in decision making and it has become an essential component in various engineering education institutions. (Mustafa Yağcı, 2022) This research is based on four steps:

1. Literature study: Previous related research work looked at understanding the importance of assessment analysis in the teaching and learning process adapted by engineering education institutions.

2. Designing and planning: Various data classification methods are studied to design and plan for machine learning training and testing algorithms to extract knowledge from large repositories for better decision-making. (Mohammed H. R. Khala et al., 2023)
3. Implementation: Preprocessing of the dataset and training the ML models using the WEKA platform by applying different performance metrics measuring assessment attributes.
4. Comparing and Analyzing: Finally evaluating the models by generating metric values and graphs.

II. MACHINE LEARNING ALGORITHMS

Machine learning algorithms used in this study are briefed as:

A. Naive Bayes

The Naive Bayes algorithm is a simple probability classifier that calculates a set of probabilities by counting the frequency and combinations of values in a given data set. The algorithm is based on Bayes's theorem where the assumption is that all variables are independent considering the value of the class variable. (Saritas, et al., 2019)

Bayes' theorem is used to determine the conditional probability $P(A|B) = (P(A) P(B|A)) / P(B)$ (Patil T. R et al., 2013)

where $P(A|B)$ is the conditional probability of the occurrence of event A when event B occurs,

$P(A)$ is the probability of the occurrence of A,

$P(B|A)$ is the probability of the occurrence of event B when event A occurs,

$P(B)$ is the probability of the occurrence of B.

B. Random Forest (RF)

It is a supervised learning algorithm used for classification and regression, used in predictive modeling and machine learning techniques (Breiman et al., 2001)

It gathers the results and the predictions of several decision trees to finally choose the best output which is the mode the value that appears most often in the set of decision trees results of the classes or mean prediction. The RF works by splitting the dataset into two sections (i) Training set and (ii) Test set. Then randomly select multiple samples from the training set. Next, use the decision tree for each sample which divides each selection into two daughters using the best division. Repeat the last step to vote for each prediction result and select the most voted prediction as the final result.

The main hyperactive parameters in RF are either used to increase the predictive power of the model or to make the model faster. A higher number of trees can increase the performance as well as make the predictions more stable, but unfortunately, the processing time becomes longer. The employment of a maximum number of features in addition to a minimum number of leaves is that the requested splitting of internal nodes may improve the algorithm's performance. Once the training step is realized, the trained model can be applied to a dataset that is not used for training. This procedure allows for estimating their predictions and comparing them to the expected values (Brownlee et al., 2019)

C. Decision Stump (DS)

It is a decision tree with a single layer. As opposed to a tree which has multiple layers, a stump stops after the first split. DS is usually used in population segmentation for large data. It is mostly used to help make simple YES/NO decision models for smaller data sets. DS is easy to build as compared to decision trees. (Peter Floml., 2010)

D. K Star (K*)

It is also called as K* algorithm. K* is an instance-based classifier, that is the class of a test instance is based upon the class of those training instances similar to it, as determined by some comparison function.

It differs from other instance-based learners (IBL) in that it uses an entropy-based distance function (DF). IBL arranges a case by assessing it to a database of pre-characterized models.

The crucial presumption is that comparative occasions will have comparative characterizations. The inquiry lies in how to characterize "comparative case" and "comparable grouping".

The comparing parts of an IBL are the separation work which decides how comparative two cases are, and the grouping capacity which determines how example similitudes yield a last arrangement for the new occasion.

The KStar calculation utilizes entropic measure, given the likelihood of changing an occasion into another by randomly choosing among every single imaginable change. Utilizing entropy as a meter for a case removal is exceptionally advantageous and information hypothesis helps in registering the space among the occasions.

The intricacy of a change of one case into an alternate is the separation among occurrences. (Mohamed Sakr et al., 2020)

E. Decision Tree (DT)

Decision Tree Classification is the process of building a model of classes from a set of records that contain class labels. DT Algorithm is to find out the way the attributes vector behaves for several instances.

Based on the training instances the classes for the newly generated instances are being found.

This algorithm generates the rules for the prediction of the target variable. This classification algorithm is used in the critical distribution of the data and makes this easily understandable.

(Hussain, M. (2018).

III. EXPERIMENTAL SETUP

In this study, all mentioned algorithms are applied to the data set using the WEKA platform for generating a confusion matrix having four possible results on Final_Result attribute values: pass, fail, distinction, and withdrawn. The performance matrices are used on confusion matrix observation values: True Positive (TP), False Positive (FP), True Negative (TN), and False Negative (FN).

$$\text{Accuracy} = (\text{TP} + \text{TN}) / (\text{TP} + \text{TN} + \text{FP} + \text{FN})$$

$$\text{Precision} = \text{TP} / (\text{FP} + \text{TP})$$

$$\text{Recall} = \text{TP} / (\text{TP} + \text{FN})$$

$$\text{Kappa} = \text{TP} / (\text{FP} + \text{TP})$$

$$\text{F1 Score} = (2 * \text{Precision} * \text{Recall}) / (\text{Precision} + \text{Recall})$$

(Alabbas Hayder., 2022)

The percentage split method is used for data validation. The data is split into 70:30 for training and testing purposes for the total data set records of 3857. The attributes considered for data set are: Code_Module, Code_Presentation, Id_Student, Gender, Region, Highest_Education, Imd_Band, Age_Band, Num_Of_Prev_Attempts, Studied_Credits, Disability, No_of_Clicks, Score, Final_Result. The Final_Result is used for considering the student performance.

These attributes are classified by considering the nature of data such as Demographic (1), Engagement (2), and Performance (3) are used as input to the model for assessing its effect on the assessment of students. The used machine learning algorithms are Naïve Bayes, Random Forest, Decision Stump, K Star, and Decision Tree.

IV. RESULTS ANALYSIS

The models are trained and tested on the WEKA platform for student datasets available on kaggle.com classified on demographics, engagement, and performance attributes using Naïve Bayes, Random Forest, Decision Stump, K Star, and Decision tree algorithms for performance metrics Accuracy, Precision, Recall, Kappa and F1 Score. TABLE I to TABLE V and GRAPH I to GRAPH V are given below.

TABLE I NAÏVE BAYES PERFORMANCE

Metrics	1	2	3	1+2	1+3	2+3
Accuracy	58.612	58.522	91.543	58.410	92.058	90.622
Precision	0.587	0.585	0.909	0.377	0.915	0.901
Recall	0.585	0.584	0.914	0.583	0.920	0.908
Kappa	0.0091	0.0099	0.8457	0.0065	0.857	0.8354
F1 Score	0.738	0.738	0.931	0.436	0.916	0.904

TABLE II RANDOM FOREST PERFORMANCE

Metrics	1	2	3	1+2	1+3	2+3
Accuracy	58.324	58.127	93.021	56.130	92.161	94.834
Precision	0.473	0.586	0.929	0.438	0.920	0.947
Recall	0.584	0.584	0.931	0.566	0.922	0.948
Kappa	0.008	0.011	0.878	0.003	0.862	0.908
F1 Score	0.008	0.011	0.878	0.003	0.862	0.908

TABLE III DECISION STUMP PERFORMANCE

Metrics	1	2	3	1+2	1+3	2+3
Accuracy	58.381	58.323	92.451	58.742	91.194	90.818
Precision	0.588	0.586	0.921	0.588	0.911	0.905
Recall	0.587	0.584	0.925	0.588	0.912	0.908
Kappa	0.011	0.011	0.864	0.012	0.838	0.834

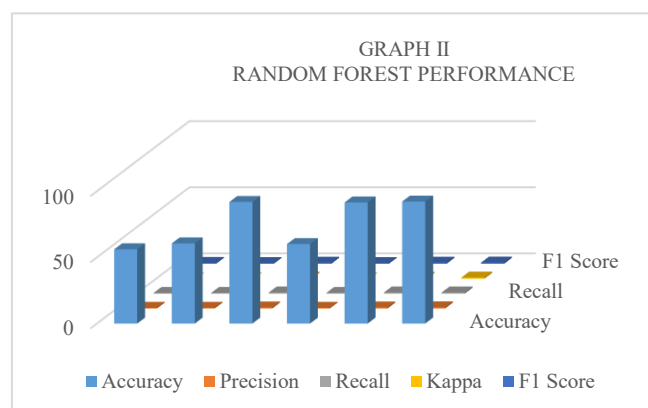
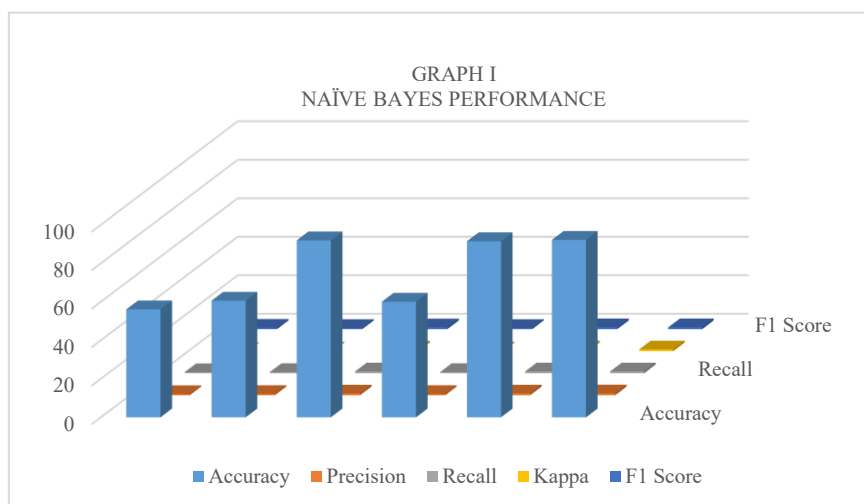
F1 Score	0.739	0.735	0.921	0.740	0.905	0.905
----------	-------	-------	-------	-------	-------	-------

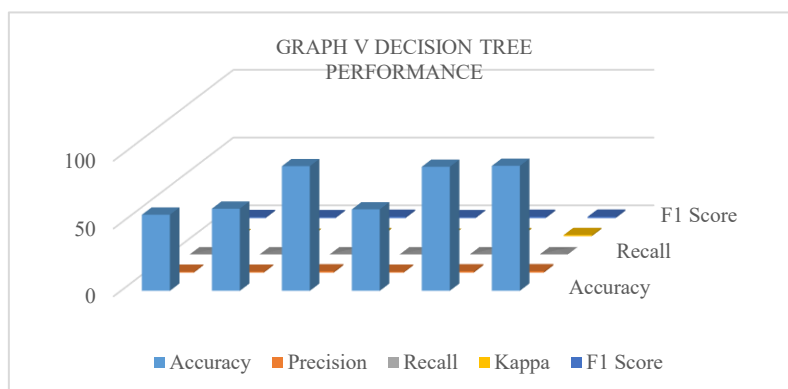
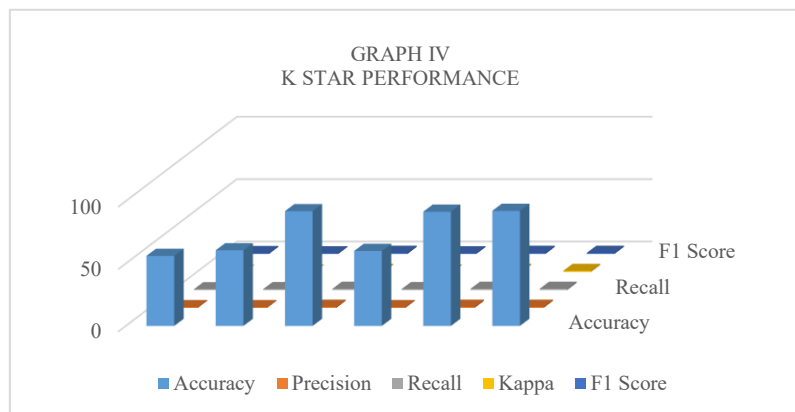
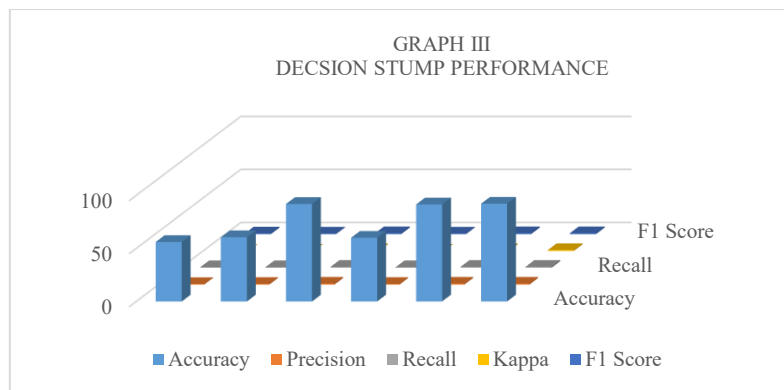
TABLE IV K STAR PERFORMANCE

Metrics	1	2	3	1+2	1+3	2+3
Accuracy	58.531	58.576	90.426	58.380	90.160	92.384
Precision	0.585	0.585	0.893	0.586	0.903	0.912
Recall	0.585	0.585	0.893	0.586	0.903	0.912
Kappa	0.000	0.000	0.825	0.001	0.817	0.862
F1 Score	0.738	0.735	0.924	0.737	0.889	0.938

TABLE V DECISION TREE PERFORMANCE

Metrics	1	2	3	1+2	1+3	2+3
Accuracy	56.124	60.604	91.944	60.078	91.516	92.215
Precision	0.590	0.640	0.907	0.612	0.91	0.927
Recall	0.561	0.607	0.914	0.601	0.915	0.927
Kappa	0.051	0.279	0.848	0.288	0.85	0.873
F1 Score	0.711	0.582	0.909	0.589	0.911	0.927





The above observation indicates that the comparison of the performance metrics shows a very mild variation in the individual and classification of the attributes identifying positive instances and maintaining stability among them.

V. CONCLUSION AND FUTURE SCOPE

WEKA platform is used for classifying student online assessment on one of the attributes into four classes: distinction, pass, fail, and dropped. Performance metrics accuracy, precision, recall, kappa, and F1 score are used to assess the classifications based on attributes such as demographics, engagement, and performance. The machine learning prediction-based model validation test is conducted using the percentage split technique. Prediction outcomes very slightly differ across the used Naïve Bayes, Random Forest, Decision Stump, K Star, and Decision Tree models and are contingent upon the chosen student classified attributes.

The WEKA platform gives the desired results, even though it can be compared with K- the fold validation technique and with some other platforms such as Rapid Miner in the next study.

This study helps in identifying the students who need special attention and allows the faculty to provide suitable counseling for a better assessment score.

REFERENCES

1. Alabbas Hayder, (2022). Predicting Student Performance Using Machine Learning: A Comparative Study Between Classification Algorithms, *Thesis Project – Mid Sweden University*
2. Anbukarasi V, A. John Martin, (2019). Student Learning Prediction Using Machine Learning Techniques *International Journal of Engineering and Advanced Technology (IJEAT) ISSN: 2249-8958 (Online), Volume-8 Issue-6, August 2019*
3. Breiman, L. (2001). Random Forests. *Machine Learning* 45(1): 5-32. . Kluwer Academic Publishers. Manufactured in The Netherlands
4. Brownlee, J. (2019). Train-Test Split for Evaluating Machine Learning Algorithms.
5. Hussain, M. (2018). Student Engagement Predictions in an e-
6. learning System and their impact on Student Course Assessment Scores. *Hindawi Computational Intelligence and Neuroscience*,
7. Mohammed H. R. Khala, and Zeinab M. Abdel Azim (2023). Predicting Student's Performance in Online Education through Deep Learning Model Predicting Student's Performance in Online Education through Deep Learning Model *Inf. Sci. Lett.* 12, No. 3, 1619-1630 (2023)
8. Mustafa Yağcı (2022). Educational data mining: prediction of students' academic performance using machine learning algorithms, *Smart Learning Environments* (2022) 9:11
9. Mohamed Sakr, Abeer Saber², Osama M. Abo-Seida and Arabi Keshk (2020). Machine Learning for Breast Cancer Classification Using K-Star Algorithm, *Appl. Math. Inf. Sci.* 14, No. 5, 855-863
10. Patil T. R., S. (2013). Performance analysis of Naive Bayes and J48 classification algorithm for data classification. *International journal of computer science and applications*, 6(2), 256-261
11. Peter Flom (2010). Building Decision Trees from Decision Stumps, SAS Global Forum
12. Saritas, M. M. (2019). Performance Analysis of ANN and Naive Bayes Classification Algorithm for Data Classification. *International Journal of Intelligent Systems and Applications in Engineering IJISAE*