



Machine Learning for Employee Promotion Analysis and Prediction

¹ Dr. P. Ramasubramaniam, ² U. Sravani,

¹ Professor, Megha Institute of Engineering & Technology for Women, Ghatkesar.

² MCA Student, Megha Institute of Engineering & Technology for Women, Ghatkesar.

Abstract

Organizations cannot function without the ability to forecast employee performance. Executives and managers who care about their companies' performance must deal with the challenging decision of promoting certain personnel since the success or failure of the business is often dependent on the competency of its employees. Most companies' present promotion policies are deceptive as they rely on managers' subjective evaluations. Using classification algorithms, this research primarily aims to build prediction models that can determine whether an employee is qualified for a promotion and, if so, which traits are most significant in determining that outcome. This article makes use of data that was submitted to Kaggle 2020. In its 54,808 rows and 13 columns, you may find details on global corporations. Organizations across nine major sectors are covered by this dataset. To forecast which employees will be promoted, we employed a number of predictive modeling approaches, such as K-Nearest Neighbors, Logistic Regression, Decision Tree, Random Forest, Support Vector Machine, and Ensemble models composed of Adaboosting and Gradient Boosting. When compared to other classification methods, Gradient Boosting performs better in terms of accuracy, F1-score, and area under the curve (AUC). Furthermore, the statistics demonstrate that an employee's rating from the previous year is the single most important factor in determining whether or not they will be promoted. Personnel advancements were unaffected by the department.

Keywords

Advancement in the workplace, AI, K-Nearest Neighbors, Logistic Regression, Decision Tree, Random Forest, SVM, Adaboosting, GTB

I. INTRODUCTION

Performance reviews are to be conducted by human resources (HR) in light of employees' contributions to the organization [1]. Crucial to the success of the business, it measures the dedication of each and every employee. It also keeps workers focused on what they should be doing for the company. Completing performance reviews by hand may be a hassle for large organizations. Employee morale and business operations may take a hit if promotions were not based on merit in such a situation. Therefore, a merit-based review and promotion procedure that is open and honest is required. A. Job Advancement Employees' careers, performance, and the company's production are all profoundly impacted by promotions, making them an essential component of any business [2]. Organizations rely on the HR department to help workers reach their professional goals. In this way, a company may amass a seasoned staff by holding on to its most productive workers; these workers can then go on to assume leadership roles with ease. An uptick in morale, loyalty, and output is a direct result of a promotion. Furthermore, promotions raise their total engagement index. Effects of Promotion (B) An employee's advancement from one level of responsibility to another is known as a promotion in the workplace. It carries with it a number of extra advantages, including a higher income, more prestige, and more responsibility [3]. Because of the increase in power, position, and authority it brings to an employee, it is a major motivator for the majority of workers. Filling openings for higher-level roles via promotions is standard operating procedure for every firm. People are more likely to put forth their best effort when they know their efforts are being appreciated. Section C: The HR manager Promotions are one of the primary



responsibilities of human resource managers. They decide whether workers have shown exceptional performance and are qualified for promotions [4]. When making important choices, the HRM usually follows the advice of supervisors from other departments. Manual suggestions, meanwhile, aren't always accurate. An employee's advancement prospects could take a hit if their supervisor is prejudiced or gives them a false report. As a result, HR is confronted with the difficult task of deciding which staff merit promotions. In addition, workers could dispute the procedure. Because a promotion brings more responsibility, more money, and maybe even a leadership position, it's crucial to assess the candidate carefully. Experience, abilities, evaluations, performance, and other criteria should be considered by the HRM before approving a promotion for an employee, and what it takes to be a leader. While some promotions are based on length of service, others take other criteria into account. But it's not always easy to objectively assess how well someone satisfies the promotion requirements. By automatically determining whether individuals should be promoted, artificial intelligence may fill this need and eliminate prejudice [5]. The purpose of this work was to use machine learning to analyze dataset attributes in order to forecast which employees would be qualified for a promotion. Employers and HR managers may also use machine learning classification models to enhance HR decision-making and the promotion process. The goals of this paper are as follows: first, to establish valid evaluation criteria for judging an employee's performance; second, to study what variables influence promotions; third, to develop a workable prediction model; and lastly, to lessen the workload of HR in finding the right candidate. This is how the paper is structured. In Section II, the literature reviews detail the relevant works. The investigation has led to a possible solution, which is detailed in Section III. Section IV, the methods used for evaluating this work. Results and comparisons of the algorithms used in the article are reported in Section V. Lastly, it shows the results and what's coming up next.

II. BACKGROUND AND RELATED WORKS

Many businesses' processes have been utterly transformed by the rapid development of information technology (IT) in recent years [6]. Information technology is crucial to the day-to-day operations of these companies' departments. An essential aspect of every company is the human resources department. Maximizing staff productivity and shielding the business from problems caused by subpar work are two of its primary benefits [6]. It connects workers with all other divisions and primarily cares for their health and safety. Human resource management mostly involves employing and dismissing workers, handling compensation and benefits, and keeping workers informed of legislation that impacts the business. They also handle staff promotions according to predetermined criteria and performance reviews. In human resource management, artificial intelligence (AI) plays a crucial role. Artificial intelligence's machine learning subfield facilitates the mechanization of analytical model construction [7]. The premise is that algorithms can learn from data, see patterns, and make judgments with little to no human input. An HR model's input values might comprise a variety of accomplishments, such as an employee's years of service and their degree of training. Therefore, a worker's performance is crucial in establishing their promotion eligibility. Workplace performance is influenced by factors such as leadership, training, and employee motivation [7]. The rate of staff turnover has also been calculated using machine learning. Experienced personnel move for better opportunities elsewhere, and the business loses a lot of them. One of the many intangibles lost when long-term workers go is the company's connection with its customers. Poor management, inadequate pay, and a hostile work environment are the main reasons employees quit [8]. Using machine learning to detect trends in employee behavior that suggest they are about to quit or switch departments is one way attrition may be controlled [9]. In order to gauge how much of an impact each employee has on a business, performance reviews are conducted. The evaluation's results are critical for informing important choices that will affect the success of the business. By supplying management with crucial information for making important choices about things like increases, promotions, and even layoffs, AI has proven to be an invaluable tool in performance reviews. By highlighting their efforts and rewarding outstanding achievement, workers find these AI exercises beneficial [10]. When workers see that the top-ranked workers are getting promotions, pay rises, and other perks, they want to be among them. Problems like prejudice or false information are common in manual performance reviews of workers. There is a risk that incompetent workers may get bonuses while dedicated ones go unrewarded. Artificial intelligence (AI) makes ensuring that performance reviews are open and honest, allowing for the advancement of top performers and the dismissal of low performers. When an individual meets management's expectations or follows established guidelines to finish their work, it usually results in a favorable performance review [10]. If an employee has been on an upward



trajectory, that fact will be shown in their performance evaluation history. An AI algorithm assesses workers' efficiency by analyzing a dataset including their work history. Keep in mind that data cleansing is essential for removing dormant employees and making sure the data is not misleading the model. When an employee completes a job to the specified criteria, we say that they have performed satisfactorily. Each employee's pay rate and potential for advancement are strongly related to how their supervisor rates their performance. As a whole, a firm benefits from the achievements of its individual members. Companies see their high-performing staff as an advantage. Recognizing and rewarding employees that continuously go above and above is our top priority. The authors Long et al. [11]. Staff Promotion Forecasting Using Job and Individual Characteristics. Based on data collected from a Chinese state-owned company, the authors of this study used machine learning techniques to forecast which employees will be promoted. The purpose of the research was to confirm that basic personal and positional data may accurately predict when an employee would be promoted. Standard classification methods such as k-nearest neighbor, logistic regression, decision tree, support vector classifier, random forest, and Adaboost were used. With an Area under the ROC (receiver operating characteristic) value of 0.96, the random forest model produced somewhat better predictions. Extra characteristics, such the amount of trainings and prizes, were not taken into account by the researchers. **Liu et al.** [12]. An Analysis of Employee Promotion Based on Data. In an effort to verify the impacts of organizational position on promotion, they used data from a Chinese state-owned business to do the study, which aimed to assess employee prospects, identify staff potential, and conduct the research. For the estimate, they used classification models such as Adaboost, logistic regression, and random forest. With an area under the curve (AUC) of 0.856, the researchers determined that the random forest model performed the best. In order to improve the promotion choices, Tang et al. [13] use usage classification and network-based approaches. In order to determine what factors may lead to a promotion for a subset of employees, this research mined information from the company's HR database. To find the best performers, we used a combination of supervised learning and graph network analysis. As part of the supervised learning process, the researchers used classification models such as Adaboost, logistic regression, and random forest. According to their findings, logistic regression outperformed the other methods with an accuracy rate of 75.61 percent. The method that yielded the greatest performance among the network-based algorithms was the one with α set to 5. Their analysis is severely limited due, in part, to the fact that they relied on a single-year dataset that lacked the leadership characteristics. To aid with employee network promotion and resignation, Yuan et al. [14] apply the regression model. Up to the year's conclusion, all 104 Strong Union workers had their social media posts analyzed. The goal of gathering this dataset was to examine the relationships between structural characteristics and workers by looking at their work-related activities and their online social networks. Employees who were given more attention on the work-related network had a higher chance of being promoted, while those who were given less attention were more likely to quit, according to the researchers who employed the logistic regression classification model. Even if their model uncovered intriguing findings, it would hold up better when compared to other models. The authors of [15] forecast employee performance using several supervised classifiers. Finding out what makes a decent employee and an outstanding one is their primary goal. One business employed machine learning to foretell how well its employees will do on the job. After data mining using the industry-wide standard procedure, the researchers used logistic regression, decision trees, and naive Bayes classification to build models for making predictions. According to the findings, logistic regression outperformed the other two classifiers in terms of accuracy. Considering the value of features might help refine the model even more.

III. PROPOSED SOLUTION

Analyzing Data Large multinational corporations (MNCs) are characterized by the dataset, which originates from Kaggle 2020 [16]. The 54,808 rows and 13 columns span nine main verticals throughout the companies. Here you may find categories such as department, area, education, gender, age, recruiting channel, number of trainings, duration of service, awards earned (yes or no), and average training score. Other information includes last year's rating. The goal characteristic is whether the promotion is yes or no. All of the factors that went into making the model are detailed in Table I.



TABLE I EMPLOYEE PROMOTION DATASET ATTRIBUTES

Features	Data Type	Description
employee_id	int64	Unique ID for employee
department	object	Department of employee
region	object	Region of employment (unordered)
education	object	Education Level
gender	object	Gender of Employee
recruitment_channel	object	Channel of recruitment for employee
no_of_trainings	int64	no of other trainings completed in previous year on soft skills, technical skills etc.
age	int64	Age of Employee
previous_year_rating	float64	Employee Rating for the previous year
length_of_service	int64	Length of service in years
awards_won?	int64	if awards won during previous year then 1 else 0
avg_training_score	int64	Average score in current training evaluations
is_promoted	int64	(Target) Recommended for promotion

Analyzing Data for Exploration Purposes

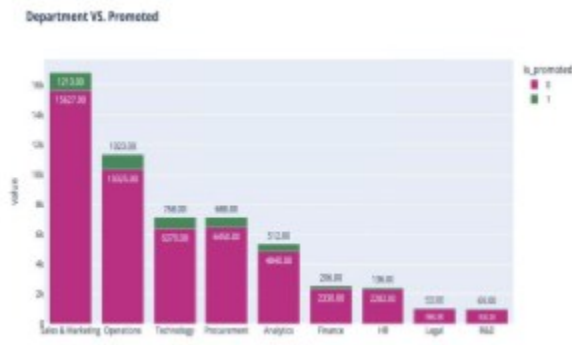


Fig. 1. Popular Departments

The majority of the company's personnel were promoted from the sales, marketing, operations, and technology departments, as shown in figure 1.

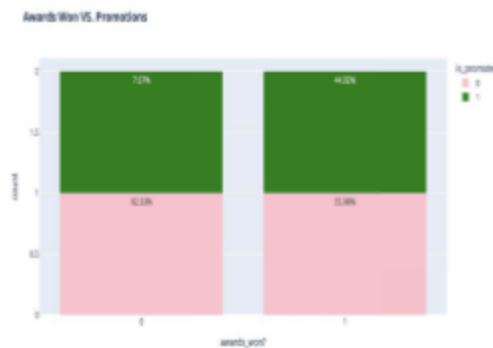


Fig. 2. Awards Won Vs. Promotion

In the figure 2 there is a high chance of getting promoted if an employee has won an award.



Fig. 3. Regions Vs. Promotion Region-2 has a disproportionately high number of promotions relative to other regions (region=7, region-22, and region-2), which is consistent with the data shown in figure 3. Two, Analyzing Correlations See how the independent variables are related in Figure 4. We made sure there was no multicollinearity by checking the correlations; multicollinearity was defined as a correlation coefficient (r) close to 0.80, which was not the case. With a weight of almost 20%, the Awards earned feature was clearly the most essential component. Age and duration of service are correlated to a degree of around 66%. Section B: Preparing Data 1) Cleaning and Preparing Data Because machine learning relies on high-quality, relevant data, data preparation is an essential step in the process.

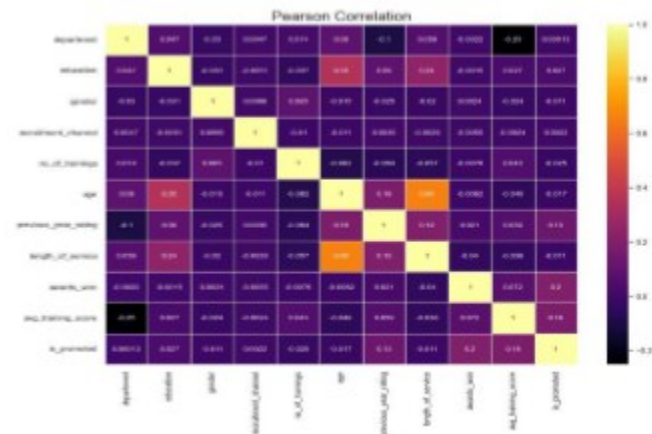


Fig. 4. Pearson Correlation

Preprocessing data before feeding it into a model is crucial since the insights gained from it directly affect the learning capacity of the models. Secondly, Data Cleaning An uneven distribution within the promoted target feature class was the first issue.



Fig. 5. Target Class Balance

To fix this, we used SMOTE, or synthetic minority oversampling. Nitesh Chawla (2002) [17] offered SMOTE, a robust approach to data imbalances. After that, we eliminated certain superfluous factors that weren't related to our analysis or predictions of the target variable, such as employee ID and area. Lastly, the dataset did not include any duplicate rows. 3) Treatment for Missing Values The education and prior year rating columns of the dataset were both missing values. Educational Aspect Since the majority of inputs were at the bachelor's level, using the mode to fill in missing values was not the best approach since it exacerbated the unbalanced data issue.

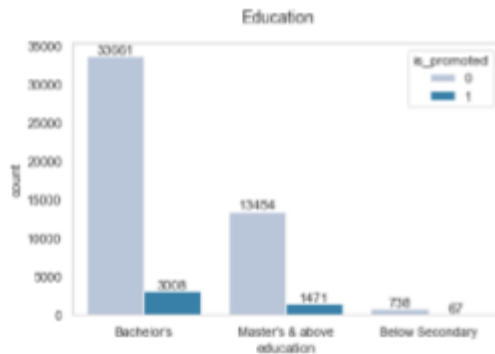


Fig. 6. Education Balance

K-Nearest Neighbors (KNN) was chosen as the method to forecast the missing data. The "previous year rating" column was set to zero by the prior Year Rating Feature. First of all, why doesn't Column "previous year rating" have any data? Data was not submitted for these workers because they were Freshers, meaning they had just been with the company for a year. The data source does not provide any information on these employees. For new hires with less than a year under their belts, we were using the figure "0" to fill in any blanks that would indicate a rating from a prior year.

4) Dealing with Outliers

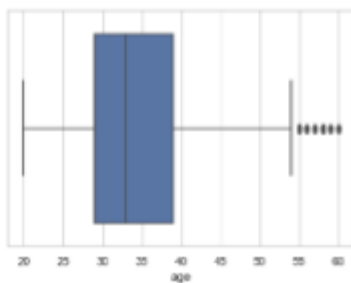


Fig. 7. Age Outlier

Outliers for the age feature between 55 and 60 are shown in Figure 5. Having said that, there are nations where the mandatory retirement age is 60 and not any younger. The age distribution was transformed into a normal distribution using three standard deviations, as age follows a normal distribution. Tests of Hypothesis The degree to which the dependent variable is reliant on the independent ones should be carefully considered [18]. It is not recommended to include an independent variable in a model's construction if it does not have a statistically significant relationship with the dependent variable. The hypothesis testing whether the two variables are associated or not was H_0 : No relationship between the two variables and H_1 : An association between the variables. The dependent variable, the target variable, was then determined by doing chi-square tests on these hypotheses. For a null hypothesis to be rejected, either (a) the p-value must be less than the significance threshold or (b) the computed test statistic must be greater than the crucial value. All of the continuous and categorical variables were significant according to the analysis of variance (ANOVA) for continuous features and the chi-square test for categorical characteristics.

6) Encoding by Category The main objective of changing the department, education, gender, and recruiting channel columns from categorical to numerical is to make them understandable to a machine learning algorithm. Categorical encoding is the process of translating categories into numerical values. One common method for encoding categorical data was utilized in this paper: label encoding [19]. This method uses alphabetical ordering to assign a unique number to each label.



department	education	gender	recruitment_channel
7	2	0	2
4	0	1	0
7	0	1	2
7	0	1	0
8	0	1	0

Fig. 8. Label Encoding

7) Scaling Data In this work, we use StandardScaler to normalize the data. By taking the standard deviation and subtracting the mean from each value, StandardScaler normalizes the values of a feature. C. Categorizing (1) Dividing Data Two sets of data were created for every model: training and testing. The model was trained using the training set, and then tested using the testing set. The data was partitioned in each model using a stratifying parameter. This ensures that the sample matches the percentage of values supplied to the parameter. At this stage, the training set received 80% of the data while the testing set received 20%. 2. Choosing the Right Model Several models were found in this study before the optimal model for the dataset was determined. To boost the model's efficiency, we tried a number of different approaches. We employed the following models: KNN, LR, Decision Tree, RF, SVM, and Ensemble (Adaboosting and Gradient Boosting), as detailed in detail below: K-nearest neighbors is an approach for data classification and prediction that does not need parameters [20]. Based on the data, the working approach estimates the classes of the vectors of the independent variables that include the most of their neighbors. By fitting data to a logistic function, logistic regression, a supervised learning approach, may estimate the likelihood of an event happening [21]. For logistic regression, the dependent variable is a yes/no binary variable where 1 denotes success and 0 denotes failure. A decision tree is a method for making decisions that employ a tree-like depiction of alternatives and potential outcomes, such utility, resource costs, and chance events [22]. This method works well with categorical datasets, which are quite common. It is easier to differentiate between promoted and non-promoted personnel using the tree-based splitting. When it comes to data classification, one supervised learning method is the Support Vector Machine (SVM) [23]. Each data point is divided into two or more categories by a line or hyperplane that runs through the middle of the dataset. One method of ensemble learning that uses trees is random forest [24]. It trains the model using a series of decision trees that randomly choose subsets of data, rather than relying on a single decision tree for data classification. In order to produce a strong classifier, the AdaBoost classifier [25] merges two weak classifier techniques. By merging many underperforming classifiers, the AdaBoost classifier is able to produce a robust and accurate classifier. One common machine learning approach for projects dealing with tabular or structured data is gradient boosting [26]. With gradient boosting, data is sorted sequentially, and new forecasts learn from the errors of older models. Using a grid search approach, the hyperparameters of the KNN and gradient boosting models were fine-tuned. A grid of hyperparameter values defines the search space, and each point in the grid is investigated via grid search. If you want to verify that a combination has worked before, this is the way to go. 3) Analyzing the Significance of Features

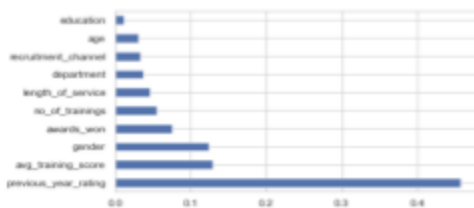


Fig. 9. Gradient Boosting Importance

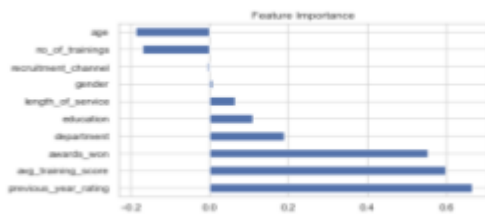


Fig. 10. Logistic Regression Importance

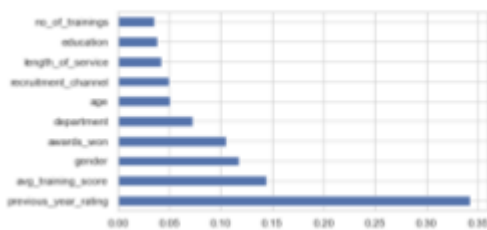


Fig. 11. Decision Tree Importance

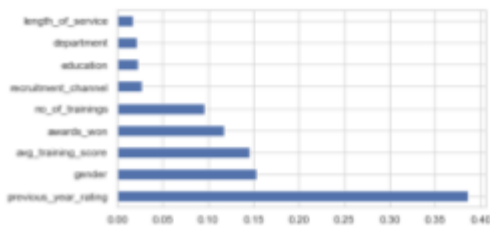


Fig. 12. Random Forest Importance

In these diagrams, The most essential factors that go into predicting a promotion are the employee's gender, average training score, previous year's rating, and awards received. Consequently, companies may highlight these crucial aspects when choosing an individual for a promotion.

IV. EVALUATION

The goal of model assessment is to identify the top performing machine learning model by comparing its strengths and weaknesses and assessing the results using a variety of evaluation criteria. Accuracy, precision, recall, F1-score, and receiver operating characteristic curve were the assessment metrics used. The accuracy and recall measures for the confusion matrix cells were determined by comparing the number of true positives and false negatives. There are four possible permutations of the expected and actual values in the confusion matrix. One well-liked statistic that integrates memory and accuracy is the F1-score, which is the harmonic mean of the two. You may think of the cell values like this: You were spot-on when you said that one of your employees would be promoted. You were right in thinking that a certain employee would not get a promotion that was suitable. You wrongly anticipated that an employee would be promoted, which is known as a false positive. You made an inaccurate prediction about an employee's promotion status, which is known as a false negative. In order to determine precision with the help of the following formula:



$$\frac{TP + TN}{TP + FP + FN + TN}$$

The following equation may be used to get the F1-score:

$$2 \times \frac{\text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}}$$

The following equation may be used to get the True Positive Rate (TPR):

$$\frac{TP}{TP + FN}$$

The following equation may be used to get the False Positive Rate (FPR):

$$\frac{FP}{FP + TN}$$

V. EXPERIMENTAL RESULTS

All of the feature selection methods were applied just before the model was fed data. The whole procedure, from data collection to modeling, relies on feature selection. The scikit-learn package comes with a number of choices for feature selection. For this research, we used this mutual information categorization strategy. To determine which independent variables provide the most useful information in regard to the dependent variable, this method computes their mutual information value. To put it simply, it evaluates the connection between characteristics and the target result. There are more dependent variables when the score is greater. An average training score of 0.030792 was determined based on mutual information, an improvement above last year's grade of 0.015075.

1) A Class II Voter The decision tree, logistic regression, random forest, and support vector machine (SVM) models were then applied using the voting classifier. A machine learning estimator, the voting classifier first trains a large number of base models or estimators before making predictions using the average of all of them. It is possible to combine the aggregating criterion with voting choices for each estimator output. In this study, we used hard voting, and the class that each classifier was most likely to predict—the class with the most votes—was defined as the projected output class. The voting classifier has an accuracy rate of 0.902663.2) Modifying the Parameters The basic idea behind a hyperparameter is that it can find the optimal parameter combination based on the scores of several preset combinations. The model scores and hyperparameter settings are shown in Table II.

TABLE II HYPER-PARAMETERS VALUES



Model	Hyper-parameter	Value	Score
K-Nearest Neighbors	n neighbors	2	0.891
	weights	uniform	
Gradient Boosting	learning rate	1	0.9395
	loss	deviance	
	n_estimators	100	
Support Vector Classifier	C	1000	0.895
	kernel	RBF	
Random Forest	max features	log2	0.885
	n_estimators	100	
	max depth	8	
Decision Tree	criterion	entropy	0.897
	min samples	8	
	splitter	random	

3) Model Testing

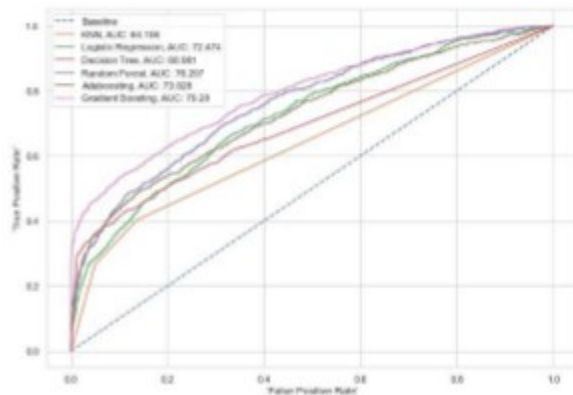


Fig. 13. ROC Curve

Results for both the training dataset and fresh samples are shown in this section, along with the methodologies that were presented. Workers who get promotions are placed in the positive class, while those who do not receive promotions are placed in the negative class. To find the area under the curve and assess the performance of each model (AUC), the receiver operating characteristic (ROC) curve was used. An analysis of each model that was applied to the dataset. When all of the models were evaluated using the area under the curve (AUC) metric, the one with the highest value—79.28—was the Gradient Boosting model.

Model	ROC Score	Accuracy Score	F1 Score
KNN:	0.62311	0.897555	0.891951
Random Forest:	0.654857	0.887247	0.889035
Logistic Regression:	0.557682	0.697409	0.762936
SVM	0.667162	0.895548	0.895775
Decision Tree	0.654356	0.899653	0.897119
Adaboosting	0.641089	0.918902	0.908412
Gradient Boosting	0.672211	0.939427	0.927597

Fig. 14. Models Performance Comparison



By calculating the accuracy, f1 score, ROC curve, and AUC—the crucial metrics for an overall evaluation—we were able to identify the best classifier for predicting whether an individual would be promoted or not. Gradient boosting was the most effective approach for the given dataset; it outperformed all other models in terms of accuracy (0.939) and F1 score (0.927), which measures a classifier's capacity to identify all positive cases. Next, the Ada Boosting classifier method achieves very high prediction accuracy (0.91 in this case). After that, the Decision Tree classifier's predictions are spot on, with an accuracy of 0.899.

CONCLUSIONS AND FUTURE WORK

A supervised machine learning classification model for promotion determination was the aim of this work. To find out who was eligible for a promotion, HR data from multinational corporations was combed through. It is critical for any firm to know which workers have the potential to be promoted. Additionally, explicitly specifying which workers qualified for promotions requires more time and effort as the organization is bigger. Making a model that can spot potential promotion prospects is, therefore, a very beneficial endeavor. Ensemble (Adaboosting and Gradient Boosting) models, KNN, Decision Tree, Random Forest, Support Vector Machine, and Logistic Regression are the prediction models that were created. Gradient Boosting fared better than the other classification methods, according to the findings. There was no evidence of prejudice, and neither the features recruiting channel nor the department had any impact on the promotion outcomes. Among the characteristics, the ratings from the previous year were the most relevant. One reasonable approach to the issue at hand is to use machine learning for the purpose of generating predicted decisions. Even with very little data, the algorithms were trained to provide respectable results. Solutions might be fine-tuned with additional data. As a result, HR analytics powered by machine learning may speed up the decision-making process and save costs. In subsequent research, we will go more into other elements that are highly correlated with the promotion issue. Also, trying to figure out how to make this model more accurate in its predictions, getting it distributed to almost all Saudi Arabian companies, and figuring out how to speed up the promotion process, as well as whether a promoted employee is qualified for a higher-level position or demonstrates desirable leadership qualities.

REFERENCES

- [1] G. R. Ferris, M. R. Buckley, and G. M. Allen, "Promotion systems in organizations." Human Resource Planning, vol. 15, no. 3, 1992.
- [2] P. Khatri, S. Gupta, K. Gulati, and S. Chauhan, "Talent management in hr," Journal of management and strategy, vol. 1, no. 1, p. 39, 2010.
- [3] J. Schwarzwald, M. Koslowsky, and B. Shalit, "A field study of employees' attitudes and behaviors after promotion decisions." Journal of applied psychology, vol. 77, no. 4, p. 511, 1992.
- [4] N. Suleman, Mahyudi, Ansir, and M. Masri, "Factors influencing position promotion of civil servants in north buton district government," vol. Volume 21, pp. PP 19–33, 04 2019.
- [5] P. Cunningham, M. Cord, and S. J. Delany, "Supervised learning, in 'machine learning techniques for multimedia'," 2008.
- [6] G. Randhawa, Human resource management. Atlantic Publishers & Dist, 2007.
- [7] P. Hamet and J. Tremblay, "Artificial intelligence in medicine," Metabolism, vol. 69, pp. S36–S40, 2017.



- [8] S. S. Alduayj and K. Rajpoot, "Predicting employee attrition using machine learning," in 2018 international conference on innovations in information technology (iit). IEEE, 2018, pp. 93–98.
- [9] P. K. Jain, M. Jain, and R. Pamula, "Explaining and predicting employees' attrition: a machine learning approach," SN Applied Sciences, vol. 2, no. 4, pp. 1–11, 2020.
- [10] B. L., "Artificial intelligence and human resource," 01 2022.
- [11] Y. Long, J. Liu, M. Fang, T. Wang, and W. Jiang, "Prediction of employee promotion based on personal basic features and post features," in Proceedings of the International Conference on Data Processing and Applications, 2018, pp. 5–10.
- [12] J. Liu, T. Wang, J. Li, J. Huang, F. Yao, and R. He, "A datadriven analysis of employee promotion: the role of the position of organization," in 2019 IEEE International Conference on Systems, Man and Cybernetics (SMC). IEEE, 2019, pp. 4056–4062.
- [13] A. Tang, T. Lu, Z. Lynch, O. Schaer, and S. Adams, "Enhancing promotion decisions using classification and network-based methods," in 2020 Systems and Information Engineering Design Symposium (SIEDS). IEEE, 2020, pp. 1–6.
- [14] J. Yuan, Q.-M. Zhang, J. Gao, L. Zhang, X.-S. Wan, X.-J. Yu, and T. Zhou, "Promotion and resignation in employee networks," Physica A: Statistical Mechanics and its Applications, vol. 444, pp. 442–447, 2016.
- [15] M. G. T. Li, M. Lazo, A. K. Balan, and J. de Goma, "Employee performance prediction using different supervised classifiers."