

Analysis of Employee Turnover factors using AI models

MSc Research Project
AI for Business

Ismail Inayat
Student ID: 24144428

School of Computing
National College of Ireland

Supervisor: Anderson Simiscuka

**National College of Ireland
MSc Project Submission Sheet
School of Computing**

Student Name:Ismail Inayat.....

Student ID:24144428.....

Programme:MSc AI for Business..... **Year:**2025.....

Module: MSc (Research) Practicum Part 2

Supervisor: Anderson Simiscuka.....

Submission Due Date:11 December 2025.....

Project Title: Analysis of Employee Turnover factors using AI models.....

Word Count: ...8538..... **Page Count:**.....27.....

I hereby certify that the information contained in this (my submission) is information pertaining to research I conducted for this project. All information other than my own contribution will be fully referenced and listed in the relevant bibliography section at the rear of the project.

ALL internet material must be referenced in the bibliography section. Students are required to use the Referencing Standard specified in the report template. To use other author's written or electronic work is illegal (plagiarism) and may result in disciplinary action.

Signature:Ismail Inayat.....

Date:27 Jan 2026.....

PLEASE READ THE FOLLOWING INSTRUCTIONS AND CHECKLIST

Attach a completed copy of this sheet to each project (including multiple copies)	☐
Attach a Moodle submission receipt of the online project submission, to each project (including multiple copies).	☐
You must ensure that you retain a HARD COPY of the project, both for your own reference and in case a project is lost or mislaid. It is not sufficient to keep a copy on computer.	☐

Assignments that are submitted to the Programme Coordinator Office must be placed into the assignment box located outside the office.

Office Use Only	
Signature:	
Date:	
Penalty Applied (if applicable):	

Analysis of Employee Turnover factors using AI models

Ismail Inayat

24144428

Abstract

Turnover is a major problem for organizations since hiring and training new staff may be expensive. It is crucial to recognize and handle new hires who may decide to change employment. Turnover-influencing traits have been identified by earlier studies evaluating turnover intention. With the goal of overcoming the limitations of previous research, the current study demonstrated a turnover intention prediction technique based on machines with deep learning. A Human Resources Analytics dataset, currently accessible to the public, was utilized in this study. Four models are tested, such as Ensemble decision tree, tuned KNN, deep scaled MLP, and TabNet. The ensemble decision tree obtained 86.05% accuracy, and the tuned KNN achieved 86.39%, above all existing KNN models. Subsequently, deep scaled MLP behaves sustainably and achieved 96.68%, indicating deep learning strength for extracting representations without manual feature extraction. The TabNet integrated attention mechanisms with feature selection using a deep learning architecture. It achieved 99.57% accuracy highest among all available TabNet algorithms. The highest performance explores an attention-based architecture suited for organized tabular data. Overall results confirmed that deep learning outperformed machine learning algorithms with the IBM HR dataset. Whereas machine learning models are less explainable and simpler. The deep learning models with improved insight features are considered suitable for HR decision making.

Keywords: Employee turnover, Prediction, AI models, machine learning, deep learning.

1. Introduction

1.1. Overview

Employee turnover is the loss regarding intellectual capital among a company's workers. Employee turnover could be costly for companies since they have to replace individuals who resign. Turnover costs several thousand dollars to each worker, including separation, substitution, and placement. Several companies struggle to comprehend employee behavior because of the various reasons for quitting a company while working somewhere. The present dissertation employs prior employee turnover information to construct a model that predicts employees susceptible to risks of turnover by considering their characteristics (Yahia, Hlel and Colomo-Palacios, 2021).

1.2. Employee Turnover

Employee turnover becomes rather than merely a recurring problem, but also requires rapid treatment. The process is an expensive issue for enterprises across industries, compromising stability and resulting in large financial losses through recruiting, training, and even productivity (Al-Suraihi *et al.*, 2021). Despite today's current business world, organizations must retain qualified personnel to ensure their survival (Urme, 2023). Organizational psychology focuses primarily on turnover intention. Turnover intention relates to people's perceptions of wanting to quit their present job. While they declined to act, this suggests a proclivity for turnover. It may be defined as the extent to which an alternative to leaving the present organization, while working for another business, is investigated and examined.

Previous research regarding turnover intention found a variety of variables influencing turnover intention. These findings contributed to increasing our knowledge regarding the turnover phenomenon. Turnover intention represents the characteristic with the most significant and strong impact on turnover (Karyati *et al.*, 2025).

To simplify the evaluation of staff attrition risk, the investigation makes use using the IBM HR Analytics Employee Attrition & Productivity dataset, another asset extensively used in the information science community. The collection of data has already been used in a variety of research and contests, with several articles delving into its insights using complex visualizations, statistical analysis, and prediction models (Raza *et al.*, 2022).

This research compares different machine learning algorithms as well as a deep learning algorithm to predict turnover. The findings during this study contribute to the continuing discussion about the use of machine learning alongside deep learning within the HR field, particularly in terms of forecasting employee behavior while maximizing workforce stability.

1.3. Machine Learning

Machine learning entails using computational methods to evaluate massive datasets. Algorithms are trained to recognize patterns in collections and use those to predict subsequent samples. Such algorithms find patterns by developing rules from information, as opposed to typical systems that utilize user-defined rules to make decisions. Models based on machines can manage large, noisy datasets containing no obvious patterns to build rules. There are several types of machine learning projects, such as unsupervised, supervised, and reinforcement learning. Particularly, supervised machine training is a method used by machine learning where the model is tasked with connecting several input characteristics to a particular target feature(s), making it the focus of this thesis. Additionally, predictions could pertain to quantitative (regression) rather than categorical classification (G. Klop, 2021).

1.4. Deep learning models

Deep learning algorithms can improve HR functions such as recruiting and staff monitoring. Using frameworks like the visual rhythm estimate allows firms to monitor worker wellness using face video analysis, enabling informed decisions regarding working settings and handling stress. Developing data-driven models with three-dimensional computational adversarial networks improves simulation understanding of workplace dynamics, including team interactions along with productivity assessments (Angermann *et al.*, 2024).

1.5. Research Problem Background

The topic of employment, particularly the decline in employment among graduates, has grown into a societal concern. Companies, on the other hand, are complaining about a human resource scarcity, which is exacerbated by the regular churn of new workers with graduated. We looked at employee and graduate turnover from two different angles. Some argue that new staff loss stems from the job searches during the first job search procedure. There exists a belief that seeking improved working conditions can be pragmatic and rational (Park, Feng and Jeong, 2024).

However, failing to settle during the early phases of a profession and accumulating volatile labor market expertise can significantly affect an individual's professional growth, future pay, and working circumstances (Munasinghe and Sigman, 2004). The churn of such fresh personnel can pose several challenges for businesses. This is due to the difficulty in recovering

the expenditures incurred during personnel selection while training. Additionally, turnover is infectious. New workers have a broader viewpoint on the business due to their familiarity with the outside workforce throughout their job hunt. Employees who watch them leave positions see issues with the business and are more likely to shift employment when given the possibilities (Chung, Kim and Kim, 2022).

As a result, promptly selecting and handling new workers with a high turnover rate helps lower the expenses of employing and managing staff. Although it is difficult for businesses to assess their intended turnover for fresh workers. Individuals sometimes conceal their turnover intentions owing to fear of punishment if exposed to others. Also, whenever it is revealed, it is difficult to evaluate whether the material is accurate. Companies are having difficulty recruiting and managing fresh employees having turnover intentions. Consequently, as a result, it is important to look at techniques to swiftly identify new employees having a high likelihood of leaving.

When there is considerable turnover among beginning professionals, forecasting a dependent factor might be difficult (An *et al.*, 2022). Identifying and managing new workers who are likely to move jobs is critical. Previous research on intended turnover has helped clarify the turnover phenomenon among new employees by identifying characteristics that influence turnover intention. Although one disadvantage of these characteristics is that they have not been managed to demonstrate whether it is feasible to forecast individuals who are likely to leave employment.

Employee turnover research has a lengthy history, with the primary focus on finding major variables, including socioeconomic and psychological aspects. Usually, regression alongside ANOVA was employed to assess these aspects, though there seems to be a significant trend toward using AI approaches for more powerful predictive analytics. Whereas algorithms including machine and deep learning have demonstrated greater efficiency in analyzing complex information, previous forecasting attempts were frequently focused on characteristics such as absenteeism with tardiness. Still, these systems occasionally lacked the full characteristics necessary for more precise predictions (Bustillo, 2025).

1.6. Research Question

- How can we use AI models to identify the most significant factors contributing to employee turnover, while comparing machine and deep learning models for accuracy?
- How effective are deep learning approaches when compared to machine learning approaches for better analysis and accurate prediction of employee turnover intention?

The remainder of the dissertation is structured in this manner: The second section provides a review of the literature, including the machine, along with deep learning methodologies and techniques employed in the current research. Section 3 covers the research approach and the collection of data, Section 4 provides a detailed implementation plan, Section 5 explores evaluation and discussion, and Section 6 is the conclusion and future work.

2. Related Work

Employee turnover research has a long history in human resource management and psychology, with the primary goal of discovering major variables, including demographic, economic, and psychological factors. Typically, regression and ANOVA were employed to

analyze these characteristics, but there seems to be a noticeable movement toward using machine learning approaches for enhanced predictive analytics (Prechsl, 2025).

J48 decision trees outperformed naïve Bayes with forecasting employee departures, obtaining an accuracy reaching 82.7% through a 70% training period split, although naïve Bayes achieved 81% through the same strategy (Iparraguirre-Villanueva *et al.*, 2024). (Fallucchi *et al.*, 2020) discovered that Gaussian naïve Bayes was the most operational model for identifying attrition, with a detection accuracy of 82.5%, while a recall rate (RR) of about 0.54, while preventing false negatives.

Additionally, studies have concluded that random forest is highly accurate, particularly since feature selection has been applied to improve the input elements, with work position, including overtime, having a major impact on attrition (Chakraborty *et al.*, 2021). Furthermore, other studies that used the random forest method for staff loss prediction disclosed a probability of 85.12%, showing its capacity for effectively recognizing key factors involving monthly income, their ages, daily rate, working years, alongside annual rate as significant predictors of staff turnover. Furthermore, (Krishna and Sidharth, 2022) demonstrated that using random forest without SMOTE for either class balance resulted in an astonishing 99.36% reliability in forecasting staff retention. Feature selection while balancing strategies dramatically improved model performance, notably in resolving class imbalance.

Meanwhile, studies comparing machine learning (ML) such as logarithmic regression alongside gradient boosting, discovered that although simpler models produce reasonable baseline predictions, ensemble approaches can improve precision as well as durability by combining multiple algorithms. Logistic regression has an 81% effectiveness in predicting shift (Najafi-Zangeneh *et al.*, 2021).

Logistic regression performed better than decision trees, along with random forest, with an accuracy reaching 88.43% (Qutub *et al.*, 2021). However, (Lazzari, Alvarez and Ruggieri, 2022) showed that LightGBM beat logistic regression along with random forest when forecasting staff leaving intention using a European dataset, with an accuracy around 64.1%.

In the work of (Yahia, Hlel and Colomo-Palacios, 2021) did a thorough research was conducted in which they integrated machine learning alongside deep learning methodologies to predict loss of staff across several datasets: medium, big, along real-world. According to the researchers' findings, the voting classifier (VC) surpassed all other models, reaching 99% accuracy across the real-world collection, 96% on the big dataset, along 98% on the medium dataset, making it a particularly effective approach to data analysis in HR.

In another research, deep learning algorithms beat standard machine learning approaches, attaining over 94% efficiency and F1 ratings, indicating their capacity to identify complex patterns inside big datasets that conventional methods may miss (Arqawi *et al.*, 2022). The study additionally found the random forest, which was the best-performing machine learning approach, with an accuracy of 92.55%. Confirmed with (Reddy, 2023), who discovered that models that use transformers outperformed classic ML models, including random forest as well as decision tree, when forecasting employee attrition, demonstrating their ability in dealing with structured and unbalanced datasets.

The work of (Raza *et al.*, 2022) stressed the importance of monthly revenue, hourly rate, work level, and aging in determining employee turnover. Their study used numerous machine learning algorithms and reached an excellent accuracy of nearly 93% for forecasting attrition,

exceeding previous cutting-edge models and demonstrating its usefulness in identifying critical factors contributing towards staff turnover.

Table 2.1 gives a complete review of significant research that has used machine learning models to predict employee turnover, emphasizing the best-performing methods and their accompanying accuracy metrics within many datasets.

Table 2. 1: Overview of studies

Reference	Dataset	Models	Top model	Accuracy
(Fallucchi <i>et al.</i> , 2020)	IBM HR Analytics	Machine learning	Gaussian naïve bayes	82.50%
(Qutub <i>et al.</i> , 2021)	IBM HR Analytics	Machine learning	Logistic regression	88.43%
(Raza <i>et al.</i> , 2022)	IBM HR Analytics	Machine learning	Extra trees classifier	93%
(Arqawi <i>et al.</i> , 2022)	IBM HR Analytics	Machine learning Deep learning	Deep learning model Random forest	94.52% 92.55%
(Lazzari, Alvarez and Ruggieri, 2022)	Europe-wide employee survey dataset	Machine learning	LightGBM	64.10%
(Bustillo, 2025)	Employee turnover in private associations from 2015 to 2023	GAs with Naïve Bayes besides KNN algorithms	FMSAX + KNN	95.52%

3. Research methodology

This research forecasted employee turnover intention using AI models such as machine learning and deep learning based approaches. The IBM HR Analytics dataset is used, which is publicly available.

3.1. Proposed solution

The proposed solution includes data gathering, experiment design pre-processing, model training, and model evaluation.

3.2. Data gathering

The dataset is selected, such as the IBM HR Analytics dataset, which is publicly available. This dataset is commonly used in various current studies (Fallucchi *et al.*, 2020; Qutub *et al.*, 2021; Raza *et al.*, 2022; and Arqawi *et al.*, 2022). This dataset is considered ideal and balanced for employee turnover prediction. The dataset is comprehensive. This research uses a dataset provided by Kaggle.com, with 35 columns (Maloku, 2024). The research presented here focuses on major characteristics that determine employee attrition statistics.

The following important variables have to be weighed:

Age: A numerical variable reflecting the age about personnel.

Attrition: The categorical variable having two options: "Yes" either "No".

Qualifications: Categorical variable including possibilities of "Below College", "High School," "Bachelor", "Master", plus "Doctor".

Staff Satisfaction: The categorical variable having possibilities of "Low", "Moderate", "High", alongside "Very High".

Work Engagement: A classification variable indicating employee involvement having options being "Low", "Moderate", "High", alongside "Very High".

Employment Life Balance: The categorical variable is that rates work-life balance whenever "Bad," "Fair," "Better," or "The Best."

Performance Rating: The categorical variable has possibilities of "Low", "Good", "Exceptional", and "Outstanding".

This data set contains additional characteristics that shall be analyzed to identify key factors affecting employee attrition. The study attempts to discover and explore relevant factors to advise methods for lowering attrition and improving organizational retention.

3.3. Experiment flow

The experiment was structured as follows. The experimental flow diagram is shown in Figure 3.1.

3.3.1. Importing necessary libraries

The first step included selecting and importing the necessary Python libraries. The dataset is in CSV format; therefore, a library is required that can manage tabular data input/output manipulation as well as data frame operations. A library is required for array operations and vector operations. Similarly, for figures and line plotting purposes library is required. The dataset analysis suggested that statistical plots are also required for this study. For dataset splitting into train and test sets library is imported. The dataset has various categorical features; therefore, a library is required for label encoding. A preprocessing unit is also required for machine learning and deep learning models. For the evaluation ROC curve, confusion matrix, and precision-recall curves are selected, and the relevant libraries are imported. Classification results, such as precision, recall, and F1-score, are required. The selected models for machine learning are KNN and decision tree, and for deep learning are Multi-layer perceptron (MLP), and Tabnet are selected. The libraries for said models are also imported.

3.3.2. Loading dataset

After selecting and importing relevant libraries, the first step is to load the dataset. The selected dataset is a CSV file, which means data is stored in a row and column format. Loading a CSV dataset means reading the dataset contents for cleaning, analysis, and training purposes.

3.3.3. Data preprocessing

To prepare the dataset for training, preprocessing is mandatory. Since most of the attributes are categorical and have to be converted into numbers. Label encoding is applied for categorical columns into numbers. In the next step data is split into training and testing sets. The training set is 80%, and the testing set is 20%. The next step is to standardize features or the scaling of features. This helped models for better performance. Since models perform better when features are scaled to a similar standardization.

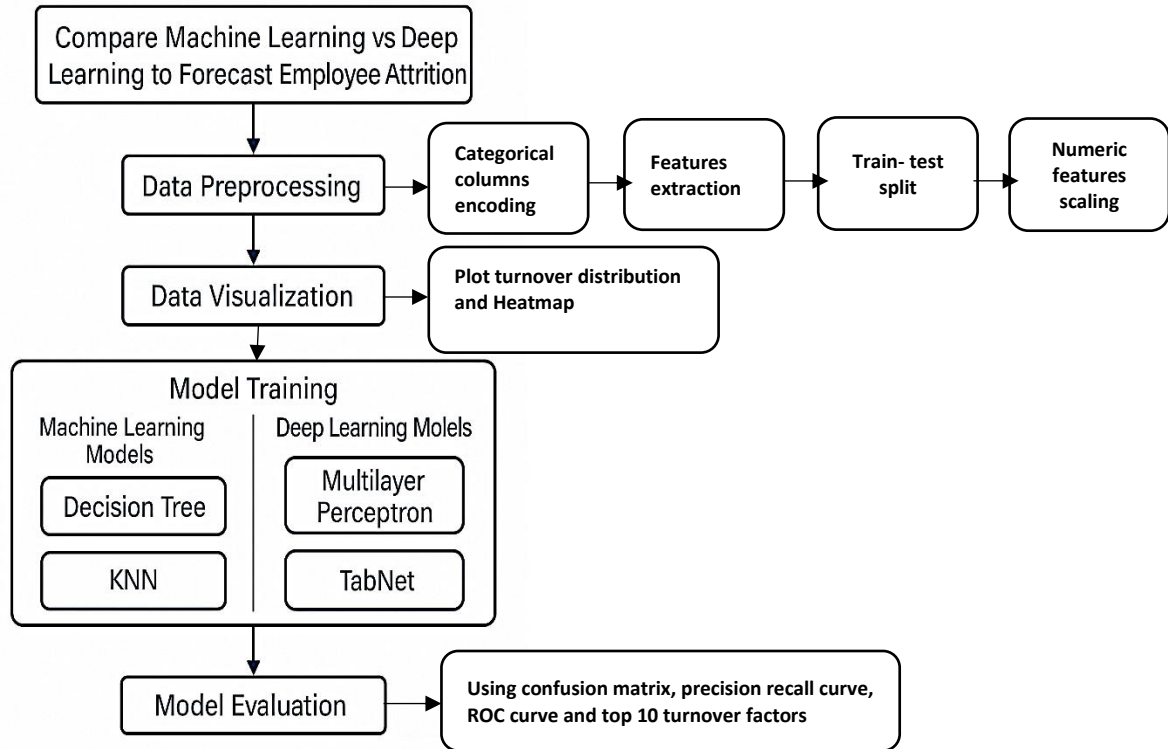


Figure 3. 1: Experimental flow diagram

3.3.4. Exploratory Data Analysis

In the next step, exploratory data analysis is performed for visualization of data to provide a better understanding before model training. A bar chart is applied to show whether employees left or stayed. The total left and stayed are counted. This process is performed for the observance of balanced data or imbalanced data. A data frame is made at this level. It contained numeric values or features. The next heat map was drawn to differentiate numeric features. Red is showing strong while blue is showing weak correlations.

3.3.5. Models

This study selected four models for employee turnover prediction. These models are selected based on the bases of literature. These models are either less explored or have low accuracy. The decision tree and KNN are trained as machine learning models for meaningful outcomes. Multilayer perceptron and Tabnet are trained as Deep learning models for more exploration of data.

3.3.6. KNN

KNN is defined and configured for training. KNN is good for five nearest neighbors for better prediction. KNN requires scaled data and is feasible for small data.

3.3.7. Decision tree

A decision tree is defined and configured for training. A decision tree splits data into branches according to features. It works as a decision process. It is simple and explainable.

3.3.8. Multilayer perceptron

MLP is a neural network-based model. Next, MLP is defined and configured for training. It is suggested for complex patterns.

3.3.9. Tabnet

Tabnet is considered a model for tabular data. The Tabnet deep learning model is defined and configured. It is better for structured data.

3.3.10. Results

In this step, each model is trained and tested respectively. Model training resulting in a prediction. Turnover predictions are compared here for correct or incorrect predictions. Next accuracy of each model is measured. Performance metrics are as follows.

Precision: true predicted left employees are correct

Recall: total of actual left employees traced

F1-score: ratio between precision and recall

Support: sample per class

3.3.11. Evaluation Matrices

The confusion matrix is designed to explore correctly or incorrectly employee's prediction. The ROC curve is also included to represent probabilities. It shows model differentiation between left employees and stayed employees. The Higher curve shows better performance of the model. It calculates False positive rates and true positive rates. Next, we plotted precision recall curve. It shows model performance on positive classes, such as left employees. It represents balance in precision and recall. Training and loss curve is plotted to record model learning. It represents model training and loss over each epoch. Whether the model is training well or not.

3.3.12. Top 10 Turnover Factor

In this phase, turnover factors are identified. Ten factors are displayed, which are most significant for turnover prediction. All models are compared for the best results.

3.3.13. Tool

All experiments were conducted on Jupyter Notebook using Kaggle. Kaggle is an online platform for Python and Jupyter Notebook. Kaggle provides 12Gb of RAM and GPU availability. Kaggle is easy to use and flexible for coding purposes.

4. Implementation

The specific experiment, working, and parameter adjustments are explained here. The experiment included importing necessary libraries, loading the dataset, preprocessing, exploratory data analysis, defining a helper function, designing, training, and evaluating Decision Tree, KNN, MLP, and Tabnet, including the top 10 turnover factors for each model.

4.1. Loading dataset

Loading a CSV dataset means reading the dataset contents for cleaning, analysis, and training purposes. The Pandas library is utilized as the dataset is in tabular form. The `read_csv()` function is applied for reading data from a CSV file. The dataset is converted into a Data Frame as a structure. In the next step, five rows from the dataset are displayed using the `head()` function for loading confirmation.

4.2. Data preprocessing

Since most of the attributes are categorical and have to be converted into numbers. Label encoding is applied for categorical columns into numbers. The function `Label Encoder()` is utilized to convert text values to numeric values. The encoding confirmation is performed to

find all text data and convert it into numeric values. This is required to convert text into numeric values, as machine learning as well as deep learning models cannot understand text data directly. The targeted variable is turnover. Therefore, features and target variables are split. In the next step data is split into training and testing sets. The training set is 80%, and the testing set is 20%. The next step is to standardize features or the scaling of features. The function Standard Scaler () is utilized for better performance and faster convergence. Scaling parameters learned using the fit_transform() function during training, and for test data scaling transform() function is applied. Since models perform better when features are scaled to a similar standardization. In the next step data frame is converted into NumPy arrays with suitable data types and labels. Furthermore, total training samples, features, data types, and labels are displayed for training understanding. The output is as under following Table 4.2.1.

Table 4.2.1: Training sample, labels and datatypes

Shapes	Data types
X_train: (1176, 34)	float32
y_train unique: [0 1]	int64

4.3. Exploratory Data Analysis

Exploratory data analysis (EDA) is a visualization of data useful for better training. A Bar chart was developed to show whether employees left or stayed using the standard plotting function. The total left and stayed are counted. This process is performed for the observance of balanced data or imbalanced data. A data frame is made at this level using the data frame function. It converted features into a labeled form. The target variable ‘turnover’ is removed at this point since it is not a feature. It contained numeric values or features. The next heat map was drawn to differentiate numeric features. It is confirmed that there are more than two numeric features; otherwise heat map doesn’t make sense. A pair correlation coefficient computed among numeric constants. Red is showing strong while blue is showing weak correlations, whereas white represents near-zero or no correlation. A heat map represents feature relations. It helped in understanding Multicollinearity.

4.4. Helper function

A helper function is designed to evaluate machine learning and deep learning models. It represented visual performance for accuracy, classification, ROC curve, confusion matrix, and precision-recall curve. An evaluation function, eval, is created. Its parameters included model name, testing true labels, models' predicted labels, and predicted probabilities for ROC and precision recall curves. Similarly, accuracies and classification reports for each model are provided. The report included each model's accuracy, F1-score, precision, recall, and support.

4.4.1. Confusion matrix

A confusion matrix is created to observe model mistakes. It helps to understand the misclassification of workers or data. It includes four attributes. True positive, False positive, False negatives, and True negative.

4.4.2. ROC curve

ROC curve created for each model. It calculated the False Positive Rate as well as the True Positive Rate using various probability thresholds. Area under the curve is computed for justification. If its score is close to 1, it is considered better.

4.4.3. Precision recall curve

It is a tradeoff for precision and recall. Precision is the correct predictive positive. Recall is the actual positives. The graph helps understand model performance in imbalanced classes.

4.5. Models

This study selected four models for employee turnover prediction. These models are selected based on the bases of literature. These models are either less explored or have low accuracy. The decision tree and KNN are trained as machine learning models for meaningful outcomes. Multilayer perceptron and Tabnet are trained as Deep learning models for exploration of data.

4.5.1. Ensemble Decision Tree

An ensemble tree classifier is designed using Random Forest, Extra Trees, Gradient Boosting, Ada-boosting, and XGBoost. It is named the pickle-safe model. This algorithm is designed to achieve the highest accuracy, specifically for employee turnover prediction. The classifiers are combined using voting. The results are evaluated, and a confusion matrix, ROC curve, and precision-recall curves are plotted. The dataset is split into 80% and 20% ratios for training and testing purposes, respectively. Stratify is used to keep a ratio of 1 for left and 0 for stay. Reproducibility is ensured with random state configuration. A 400 decision trees were built with random subsets of features. Class imbalance is handled with a class weight function. An extra classifier is developed for extra randomness for the splitting tree. Adaboost classifier is applied for boosting weak learners towards a strong model. XGBoost is a powerful boosting model that has speed along with accuracy optimization. For better generalization, randomness was added using colsample and subsample. Multithreading is disabled to ensure for the saving model. Each algorithm predicted a probability for classes. Soft voting combined the average probability of the model with steady and robust performance instead single decision tree. All algorithms and parameters are serialized and saved, avoiding the issue of parallel processing. All classifiers are combined and trained. Predicted labels and predicted probability are calculated. Accuracy, ROC-AUC calculated. Confusion matrix and the ROC curve are plotted.

4.5.2. Tuned KNN

The optimized and fine-tuned KNN is designed and implemented for this case. This algorithm is scaled, tuned, and bagged to achieve high accuracy. A pipeline of scalar and KNN has been designed. Scalar standardized numeric values and KKN utilized a classifier. The data is automatically scaled whenever training and testing are performed. A grid of hyperparameters has been set. A grid search was set up for testing uniform nearest neighbors. The closer neighbor is considered more influential. Two types of distance, Euclidean and Manhattan, are applied. A cosine metric is set to compare similarity among directions instead of distances only. A grid-based search is designed to achieve the best setting for KNN. This grid search trained several versions for the pipeline, utilizing various combinations of hyperparameters. A reliable evaluation was applied using cross-validation by 34 cross-validation. The best model is retrieved for KNN as fine-tuned and optimized. A bagging classifier is used to wrap the best KNN using an ensemble having 249 various slightly different models. The training sample of 90% seen by KNN. Each KNN utilized 90% features. All these steps reduced overfitting as well as generalization improvement. Next, KNN training is performed with a bagged ensemble with the training set for predicted classes and positive classes probabilities. Accuracy for training data was evaluated to check for overfitting problems. Last but not least, evaluations

and plots such as the confusion matrix, precision, recall, and ROC curve were produced.

4.5.3. Deep Scaled Multilayer Perceptron

An enhanced deep MLP is designed. Numerical values are normalized with a scalar function for input features. MLP performs well when the data is normalized. Similarly, scaling confirms an equal contribution of all features. Deep MLP is configured with 2 hidden layers having 2048 and 1024 neurons to design a deep network. An efficient deep learning activation function, ReLU, is applied. For adaptive learning Adam optimization algorithm is included.

Overfitting is prevented through L2 regularization with $1e^{-3}$. Learning rate is set constant for fixed training. For stable and slow learning, initial learning is set to $3e^{-1}$. A large number of epochs for training is used. Early stoppage of training is applied while validation stops improving. It is set to 20 as consecutive improvements are not achieved. Reproducibility is ensured. Mini-batch sample of 12 is adjusted. Training data of 10% is reserved for early collapse. A neural network is trained on scaled data. Internal biases when minimizing classification loss. Prediction and probability are collected. Evaluation is performed using accuracy, confusion matrix, ROC curve, and precision-recall curve.

4.5.4. Tabnet

Model's hyperparameters are adjusted to 32 No. of decisions, and 32 are attention features to control the model's capacity. Decision steps are adjusted to 5. Features reusability is adjusted to 0.1. Sparsity is $1e^{-3}$ to avoid overfitting. The adaptive learning is achieved using Adam. Learning rate is adaptive and set to 0.2. For the feature selection procedure, sparsimax is chosen, which chooses only relevant features. Learning rate schedule is set to reduce it for every epoch. LR is reduced by multiplying gamma after step epochs.

Reproducibility is ensured by a random seed. After the adjustment of the hyperparameters, model training is started. Training performance as well as validation performance is evaluated after every epoch. The training process stops if improvement is not increased after 50 epochs. A small and stable batch size is chosen. The final batch is kept even if it is smaller. Overfitting and underfitting are observed. Predictions and probabilities are observed. The probability of an employee leaving class is extracted. The model is evaluated by calculating accuracy and plotting a confusion matrix, precision, recall curve, and ROC curve.

4.6. Predicting the top 10 turnover features

The top ten turnover features are predicted for each model. First feature importance for TabNet is collected, and a bar chart is plotted. The next top ten turnover features for the Ensemble decision tree are selected. Turnover predicted features are collected from the voting classifier, as all tree algorithms are ensembles in the voting classifier. The predicted features are collected, and a bar chart is plotted. The predicted features for KNN are collected through permutation importance for estimation since KNN does not have feature importance as built-in. The predicted features are collected, and a bar chart is plotted. Finally, the top ten predicted features of the deep ensemble MLP are collected. A SHAP is used for a kernel explainer. A mean SHAP is computed for features. The predicted features are collected, and a bar chart is plotted.

5. Evaluation

This section's goal is to give a thorough examination of the research's primary findings and outcomes, along with their implications from an academic and practical standpoint. The study question and objectives are supported by just the most pertinent findings. The results are thoroughly and rigorously analyzed. The results of the experimental study are evaluated critically, and their predictive levels are determined.

5.1. Exploratory Analysis

After importing the dataset and performing preprocessing, the first test is made to visualize and explore the dataset. The bar chart in Figure 5.1 shows the employee turnover distribution. Almost 1200 employees stayed, and 250 are left.

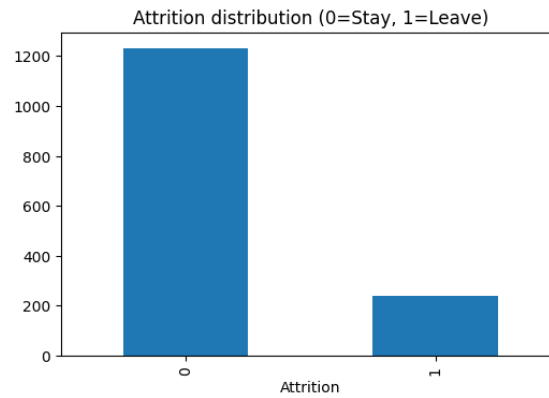


Figure 5. 1: Turnover Distribution

In the next step, a correlation heat map is generated as shown in Figure 5.2. It showed the relationship of employee features from the dataset. Positive and strong correlations are shown in red. Negative strong correlations are in blue. No correlations or weak correlations are in light colors. The diagonal represents in red as each variable is strongly correlated with itself. This heat map is generated randomly on a dataset without the proposed models' implementation. The positively strong correlation variables are years at company, years with current manager, years in current role, and years since last promotion.

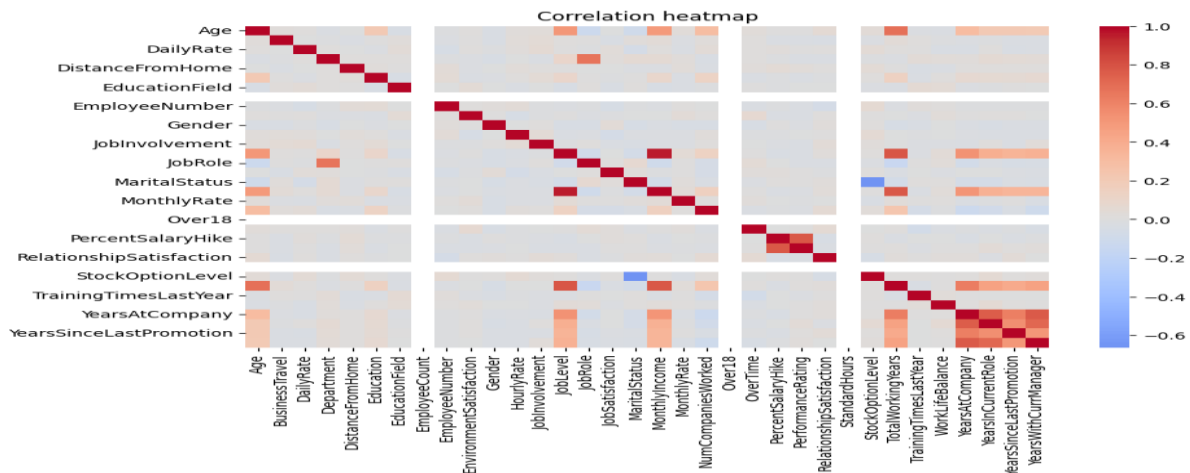


Figure 5. 2: Heat map

It means an employee who has had a longer time has a positive relationship with the manager and a long time in the current role. Age shows a positive moderate correlation with years at the company. There is no correlation or no correlations among distance from home, gender, and educational area with different variables. The job role and stock level have a weak relation

through other features. The heat map is useful in this work to identify Multicollinearity. Since some variables are very similar, such as years in the company and years with the manager.

5.2. Ensemble Decision Tree

After training and implanting the Ensemble decision tree first confusion matrix is plotted for evaluation. The true negatives are 242; here means stayed employees correctly predicted to stay. The false positives are 5 employees predicted as having left. False negatives are 36 is showing left employees but predicted as stayed. True positives are 11 left employees, but predicted to leave. The proposed ensemble decision tree performed well for the stayed class as 242 from 247 predictions being correct. Although it behaves poorly for the leavers class, as 11 out of 47 predictions are correct. The overall accuracy achieved is 86.05% with 69% precision and 23 % recall. Recall for left employees is very low, indicating the model missed many turnover cases. The confusion matrix is shown in Figure 5.3.

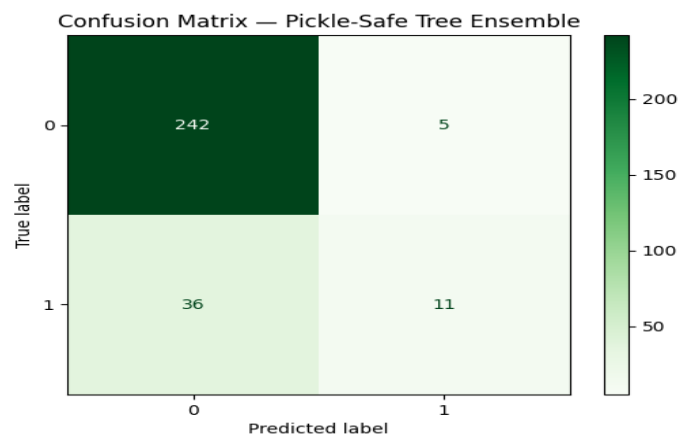


Figure 5. 3: Confusion Matrix for ensemble decision tree

The ROC curve is also designed and plotted for cross-evaluation and justification. The ROC curve is especially designed for classification performance. The blue line is indicating definite ROC curve of the ensemble decision tree. Gray line indicates a random classifier with an AUC of 0.5. The area under the curve is 80.90%. It shows the model is good ability of discrimination. The Model has 80.90% chances for positive instances than negative instances. The 80.90% suggesting a good classification performance. The ROC curve for the ensemble decision tree is shown in Figure 5.4.

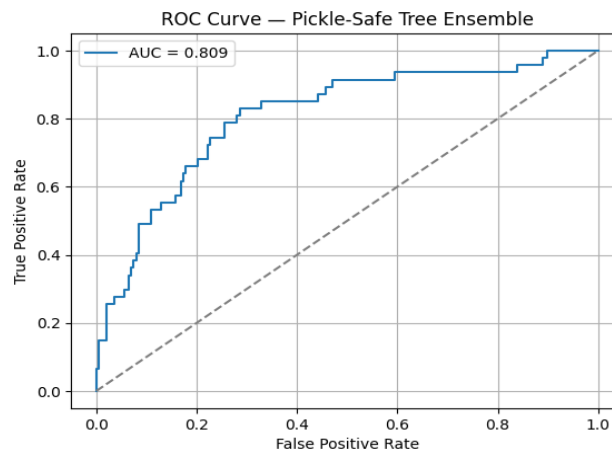


Figure 5. 4: ROC curve for ensemble decision tree

5.3. Tuned KNN

This algorithm is scaled, tuned, and bagged to achieve high accuracy. A pipeline of scalar and KNN has designed. Scalar standardized numeric values and KKN utilized a classifier. The traditional KNN is not best fitted with the dataset. Therefore, tuned KNN is designed and implemented for better turnover prediction performance. The tuned KNN achieved 86.39% accuracy with 86.83% precision for class 0 and 76.92% for class 1. The recall is 98.79% for class 0 and 21.28% for class 1. Similarly, the F1 score is 92.42% for class 0 and 33.33% for class 1. The confusion matrix is designed and plotted for better evaluations and justifications. The model behaves excellently as its true negative is 244 and false positive is 3 for class 0. Although the model misses class 1, as its false negative rate is 37. It means the dataset is not correlated with tuned KNN. The tuned KNN performed conservatively with 0 classes. The confusion matrix for tuned KNN is shown in Figure 5.5.

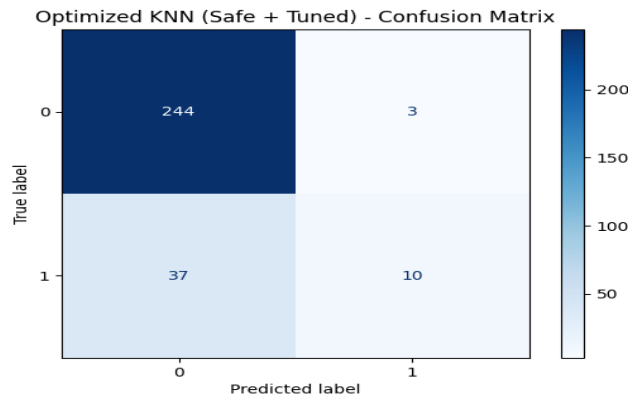


Figure 5. 5: Confusion matrix for tuned KNN

The ROC curve is also designed and plotted for tuned KNN. It helps with classification performance. The area under the curve is 74.8%. It indicated model quality is good or reasonable. The tuned KNN correctly ranked positive instances. The ROC curve does not achieved high AUC, indicating that it misses class 1. The ROC curve for tuned KNN is shown in Figure 5.6.

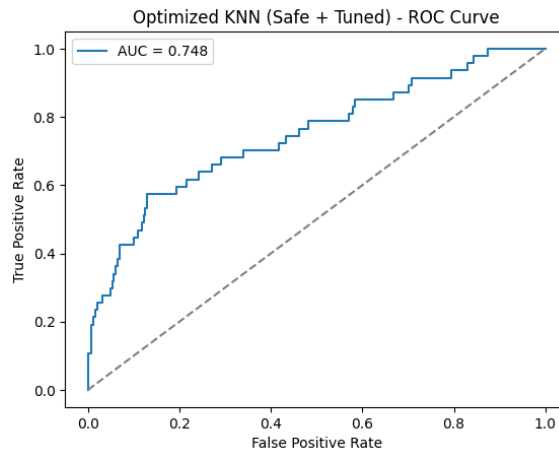


Figure 5. 6: ROC curve for tuned KNN

Another justification is made with precision recall curve. The curve represented that how many positive actual samples were correctly identified, named Recall on the x-axis. Similarly, how many positive predicted samples were actually positive, named Precision on the y-axis. The curve is showing that whenever recall was low, the precision seemed high. When recall

increases for positives, precision drops. The model struggled while maintaining precision as it started getting positive class labels. The precision-recall curve is shown in Figure 5.7.

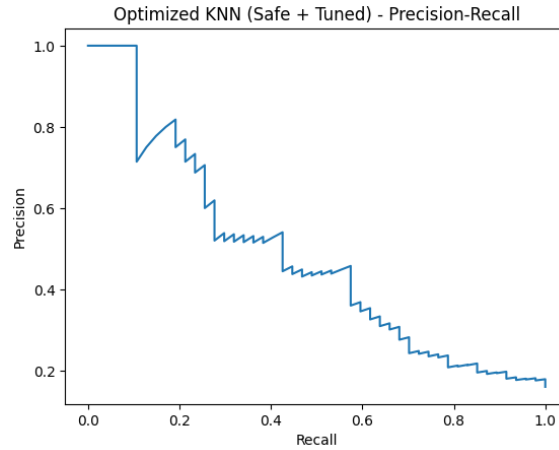


Figure 5. 7: Precision-Recall for tuned KNN

5.4. Deep Scaled Multilayer Perceptron

An enhanced deep MLP is designed. Numerical values are normalized with a scalar function for input features. MLP performs well when the data is normalized. Similarly, scaling confirms equal contribution of all features. It achieved ensemble accuracy up to 96.68% and the precision with class 0 is 87.74% and 45.45% for class 1. Its recall for class 0 is 92.72% and 31.91% for class 1. Similarly, it has a 90.16% F1-score for class 0 and 37.59% for class 1. The first confusion matrix is plotted for a deep scaled MLP. The true negative is 229 for predicted class 0, as it was true 0. Its false positive is 18 indicated that the predicted class is 1 when the actual class is 0. The false negative is 32 and the model failed to predict class 1. The true positive is 15 when predicted, and the true classes are 1. The model proved good for class 0 but struggled when class is 1. The confusion matrix is shown in Figure 5.8.

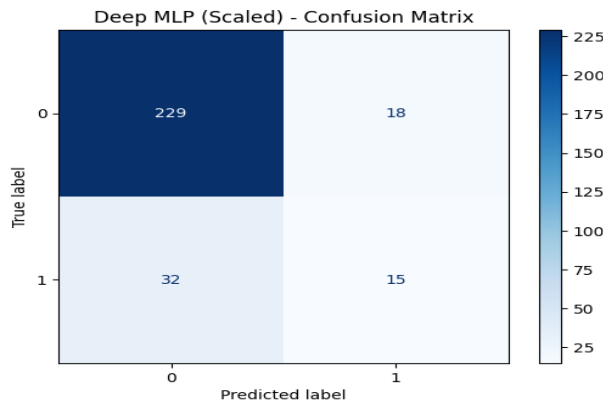


Figure 5. 8: Confusion matrix for deep scaled MLP

The ROC curve is also plotted for deep scaled MLP. The x-axis is set for the false positive rate, and the y-axis is set for the true positive rate. False positives are the ratio of negatives classified as positives, and true positives are the ratio of correctly identified positives. The AUC is 70.20% shows quantifying performance. It is moderate level, neither poor neither best. The model is still weak for the minority class. The ROC curve is shown in Figure 5.9.

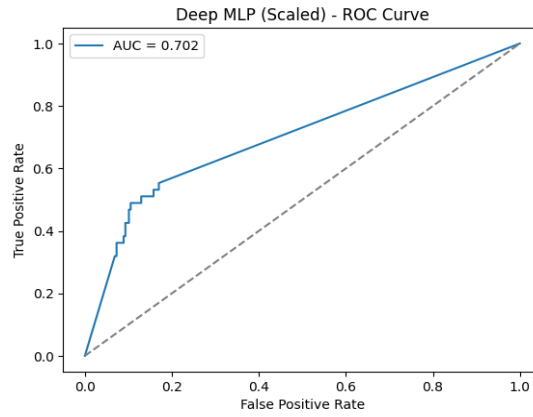


Figure 5. 9: ROC curve for deep scaled MLP

The precision-recall curve is also designed and plotted for the deep scaled MLP. The curve precision is 1, with 0 recall indicating that the model is correct during positive predictions. However, the model missed actual positives mostly with low recall. The precision-recall curve is shown in Figure 5.10.

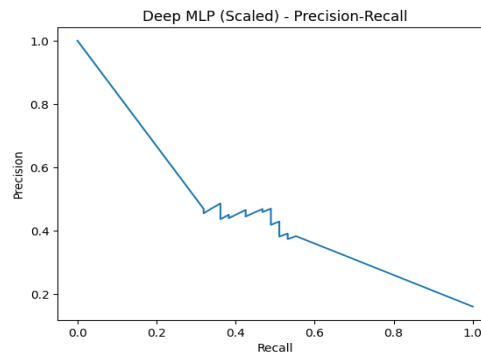


Figure 5. 10: Precision recall curve for deep scaled MLP

5.5. TabNet

The Tabnet deep learning model is defined and configured. It is better for structured data. The Probability for employee leaving class is extracted. The model is evaluated by calculating accuracy and plotting a confusion matrix, precision recall curve and, an ROC curve. The model achieved 99.57% accuracy and 90.94% precision for class 0 and 60% for class 1. Its recall is 93.52% for class 0 and 51.06% for class 1. The F1 score is 92.22% for class 0 and 55.17% for class 1.

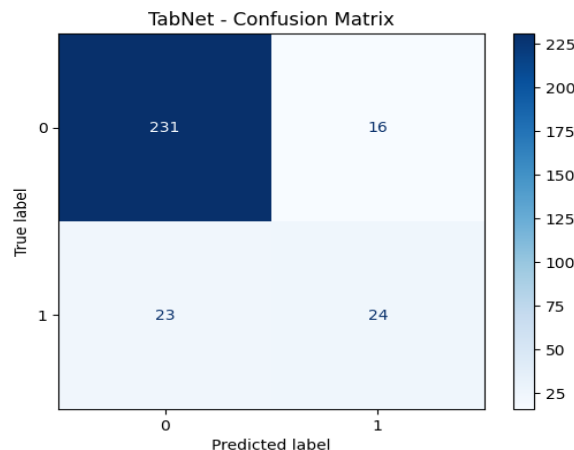


Figure 5. 11: Confusion matrix for TabNet

The confusion matrix is also plotted for this model. The true negatives are 231, which means class 0 instances are correctly predicted. The false positives are 16 here, incorrectly predicted for class 1, while the true labels were class 0. The false negatives are 23 instances predicted as class 0 as incorrect. These are missed cases. The true positives are 24 instances, indicating correctly predicted for class 1. The overall correct predictions are 255, and 39 are incorrect predictions. The confusion matrix is shown in Figure 5.11.

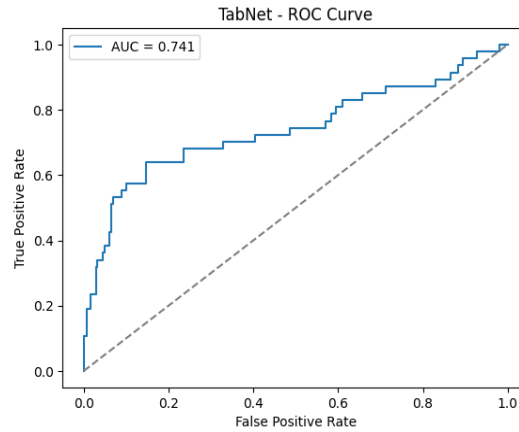


Figure 5. 12: ROC curve for TabNet

The ROC curve is also plotted for this model. The AUC is 74.1%. It is representing that the model is able to distinguish positive and negative classes instead of random, still needs improvements. The ROC curve for TabNet is showing in Figure 5.12.

The precision-recall curve is also plotted for this model. Precision on the y-axis measured the ratio of positive predictions as they were actually correct. There are low false positives if high precision. Recall is on the x-axis, known for sensitivity. It is ratio of correctly identified actual positive instances. There are low false positives if the recall is high. The curve is showing if recall increased, the precision decreased. This suggested that for more positive cases model should predict more or more false positives. The precision-recall curve for this model is shown in Figure 5.13.

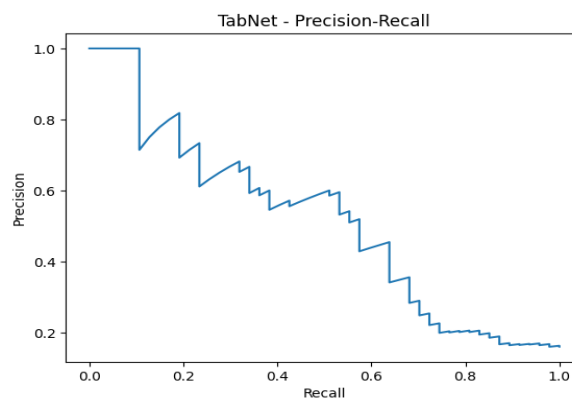


Figure 5. 13: Precision recall curve for TabNet

5.6. Predicting the top 10 turnover features

After implementing four proposed models, employee turnover prediction is made. The top ten turnover features are predicted for each model, such as ensemble decision tree, KNN, deep scaled MLP, and TabNet.

5.6.1. Predicting the top 10 turnover features using an Ensemble decision tree

First top ten features predicting using an Ensemble decision tree. The Ensemble decision tree predicted the following features that are causing employee turnover using the IBM HR dataset. The predicted features using the Ensemble decision tree are shown in Figure 5.14.

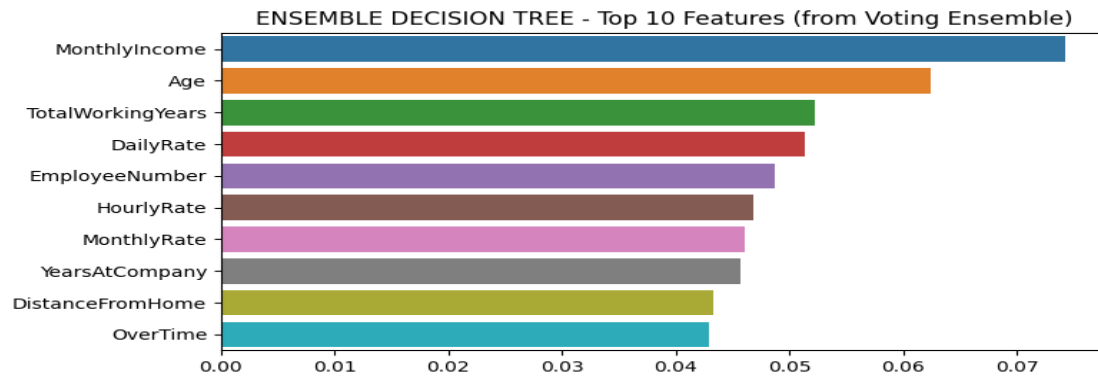


Figure 5. 14: Top 10 turnover predicted features using an ensemble decision tree

According to the Ensemble decision tree prediction, monthly income is the most significant factor that is causing employee turnover. The second factor is age that is causing turnover. Total working years, daily rate, hourly rate, monthly rate, and years at company have less influence but also causing turnover. Distance from home and over time are the least significant features that were predicted by the Ensemble decision tree caused employee turnover.

5.6.2. Predicting top 10 turnover features using tuned KNN

The predicted top 10 turnover factors by tuned KNN. The tuned KNN predicted over time is most significant factor. Job satisfaction is the second most influential factor playing a role in turnover. The third and fourth are environment satisfaction and stock option level predicted by tuned KNN. The next influence is created by training times last year, job level, marital status, and number of companies worked. The least predicted factors for turnover by tuned KNN are performance rating and distance from home are shown in Figure 5.15

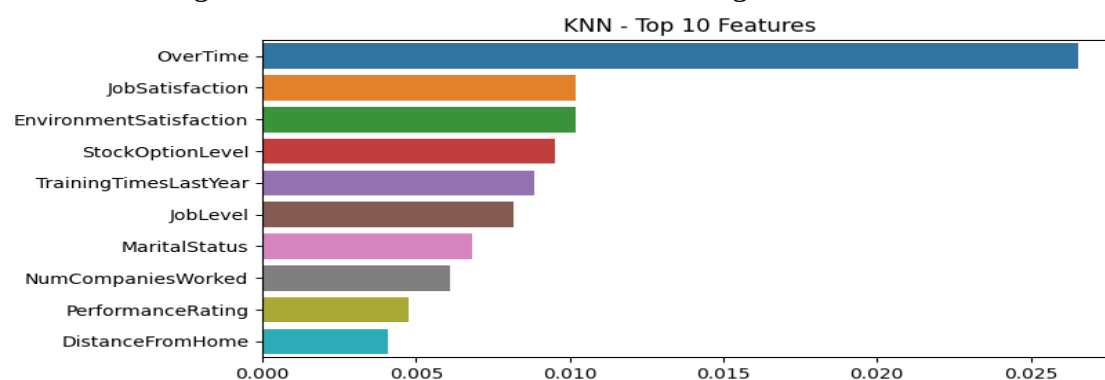


Figure 5. 15: Top 10 turnover predicted features using tuned KNN

5.6.3. Predicting top 10 turnover features using deep scaled MLP

The deep scaled MLP predicted that overtime is the highest significant factor for turnover. The Number of companies worked is ranked second. Similarly, gender at third, department at fourth, stock option level ranked at five, monthly income at six, job satisfaction is at seven, years in current role is at eight, environment satisfaction is at nine, and least significant is years

with current manager ranked at ten.

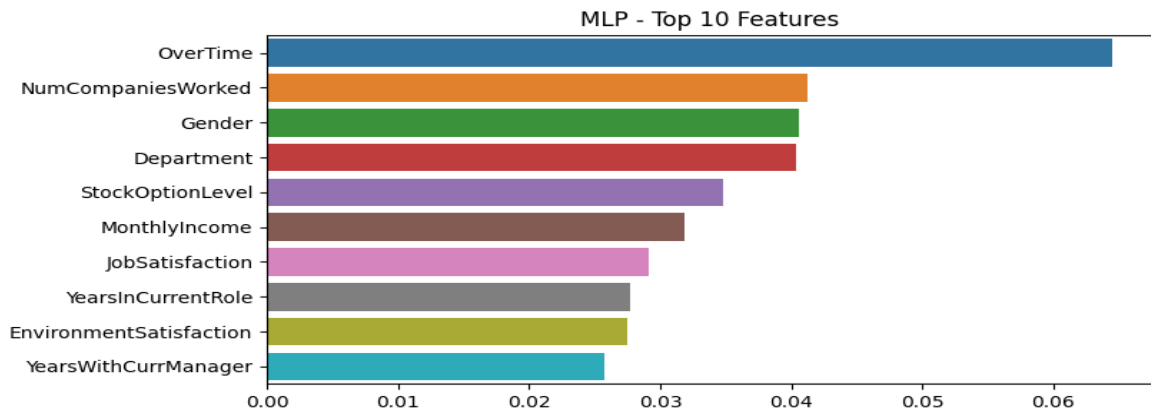


Figure 5. 16: Top 10 turnover predicted features using deep scaled MLP

Deep scaled MLP predicted experimental and behavioral factors, such as over time and the number of companies worked. It also predicted organizational context and personal satisfaction, representing a balanced prediction for employee turnover. The predicted turnover factors by using deep scaled MLP are shown in Figure 5. 16.

5.6.4. Predicting top 10 turnover features using TabNet

The department factor is predicted as the most influential factor for employee turnover using TabNet. The next factors are Training Times Last Year, number of companies worked, education level, years in current role, job satisfaction, monthly income, marital status, and age, which are ranked two to nine, respectively. The least influential factor predicted by TabNet is the education field. The TabNet predicted structural as well as organizational context. The TabNet predicted factors are shown in Figure 5.17.

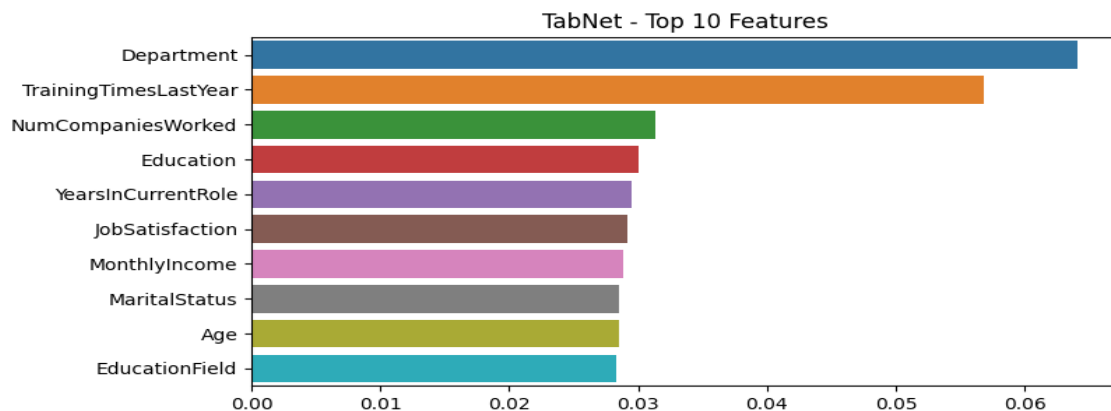


Figure 5. 17: Top 10 turnover predicted features using TabNet

5.7. Discussion

The results of the evaluations are discussed in depth in the following paragraphs. This section offers criticism of the experiment alongside an honest assessment of whether the methodology was adequate. Recent research paper results are contrasted with model evaluations.

5.7.1. Comparative study on an ensemble decision tree

The overall accuracy achieved by my ensemble decision tree model is 86.05% with 69% precision and 23 % recall. Recall for left employees are very low, indicating the model missed many turnover cases. In the work of (Živković and Šebalj, 2024) and (Ahmed, S.R., Ahmed,

A.K. and Jwmaa, 2023), authors only implemented a single decision tree, such as CART and J48. The tree is sometimes over-fitted. The Single tree has low predictive accuracy than the proposed ensemble decision tree. While in work of (Venkat, Chaitanya and Teja, 2023) applied a simple decision tree. The proposed work reduces variance and bias, improves robustness, and handles nonlinear interactions, which are the limitations of previous works. Similarly, model complexities are low, medium, and low, respectively, in said works, while this work has high model complexity. The Table 5.1 shows that the proposed ensemble decision tree attain highest accuracy over other tree algorithms with turnover datasets.

Table 5. 1: Comparison of the proposed ensemble decision tree with recent decision trees

Reference	Model	Dataset	Accuracy
(Živković and Šebalj, 2024)	Random Tree	Turnover Intention Dataset	60%
(Venkat, Chaitanya and Teja, 2023)	CART Decision Tree	IBM HR Dataset	77.8%
(Ahmed, S.R., Ahmed, A.K. and Jwmaa, 2023)	Decision Tree	IBM HR Analytics	84.2
This Work	Proposed Ensemble decision tree	IBM HR Analytics	86.05%

5.7.2. Comparative study on tuned KNN

The traditional KNN is not best fitted with the dataset. My proposed tuned KNN achieved 86.39% accuracy with 86.83% precisions for class 0 and 76.92% for class 1. The work of (BALCIOĞLU and ARTAR, 2022), (Atef, Elzanfaly and Ouf, 2022), and (Fallucchi *et al.*, 2020) performed basic preprocessing, while the proposed work used scaling and pipeline based preprocessing. Similarly, cited works used minimal and variable tuning of hyperparameters tuning, respectively, but this work used GridSearchCV for hyperparameter tuning. Cross validation is low among cited works, while this work applied a high CV. Cited works applied basic and rare ensemble method respectively, and this work applied bagged KNN as ensemble learning. This work is highly stable as compared to the cited works, as they are medium stable. Last but not least, as per the quality of engineering, cited works are low in quality, while this work has used a professional pipeline and pickle safe. The Table 5.2 represents the proposed tuned KNN achieve best accuracy comparing with recent articles.

Table 5. 2: Comparison of the proposed tuned KNN with recent KNN models

Reference	Model	Dataset	Accuracy
(BALCIOĞLU and ARTAR, 2022)	KNN	IBM HR Analytics	81.0%
(Atef, Elzanfaly and Ouf, 2022)	KNN	IBM HR Analytics	84.0%
(Fallucchi <i>et al.</i> , 2020)	KNN	IBM HR Analytics	84.2%
This Work	Proposed tuned KNN	IBM HR Analytics	86.39%

5.7.3. Comparative study on deep scaled MLP

In my work, numerical values are normalized with a scalar function for input features. MLP performs well when the data is normalized. Similarly, scaling confirms equally contribution of all features. My model achieved ensemble accuracy upto 96.68%. The work of (Alqahtani,

Almagrabi and Alharbi, 2024) used a simple decision tree, (Shiri and Yamaguchi, 2025) applied deep Bi-TCN and Durães, 2024) applied random forest, while this work applied deep MLP with fully connected. The first cited work applied deep learning, and the third applied basin ANN, while this work applied tuned hyperparameters. First cited work has no optimization of hyperparameters, 2nd used extensive tuning, and the third with light tuning, while this work has aggressive tuned with batch size, overfitting is not prevented in the first cited work, 2nd work applied dropout, and the third applied train validation concept, while this work applied L2 alpha. Primarily, the 1st work is survey based application, 2nd is a predictive temporal, and the third is employee churn prediction, while this work has high classification accuracy with deep learning. The Table 5.3 represents the proposed deep scaled MLP achieve best accuracy compared with recent articles.

Table 5. 3: Comparison of the proposed deep scaled MLP with recent MLP

Reference	Model	Dataset	Accuracy
(Alqahtani, Almagrabi and Alharbi, 2024)	Multilayer perceptron	IBM HR Analytics Employee Attrition dataset	88.43%
(Shiri and Yamaguchi, 2025)	Multilayer perceptron	IBM HR Analytics	86.93%
(Durães, 2024)	Multilayer perceptron	IBM HR Analytics	92%
This work	Proposed deep scaled MLP	IBM HR Analytics	96.68%

5.7.4. Comparative study on TabNet

In the work of (BALAZS GONCZY, 2022) and (Ramasamy, Muniandy and Subramanian, 2025) applied partial deep learning, and (Seth, 2025) applied partial deep learning as ANN, while this work attention based TabNet. All cited work applied manual feature learning, while this work applied attention-based automatic feature learning. The First cited work has high interpretability, 2nd and 3rd having medium, while this work has high interpretability. The first cited work with basic optimization, 2nd with regular ML tuning, and 3rd with grid or random search, while this work applied LR scheduling, Adam and sparse regularization. Case similarity is very high between the first two works and moderate in the third, while this work applied a direct match mechanism. Some unique features, such as the probability for employee leaving class is extracted. Automatically feature learning, sparse attention, end-to-end feature learning, advance optimizers made this model state-of-the-art to achieve high accuracy and interpretability among all. The model achieved 99.57% accuracy and 90.94% precision for class 0 and 60% for class 1. The Table 5.4 represents the proposed TabNet achieve best accuracy comparing with recent articles.

Table 5. 4: Comparison of the proposed TabNet with recent TabNet models

Reference	Model	Dataset	Accuracy
(BALAZS GONCZY, 2022)	TabNet	IBM HR Analytics	86%
(Ramasamy, Muniandy and Subramanian, 2025)	TabNet	IBM HR Analytics	89.6%
(Seth, 2025)	TabNet	IBM HR Analytics	81.55%
This work	Proposed TabNet	IBM HR Analytics	99.57%

5.8. Summary

The first machine learning model is an ensemble decision tree that obtained 86.05% accuracy and outperformed existing tree models. The second machine learning model is tuned KNN achieved 86.39% above than all existing KNN models. The results explore conventional algorithms that captured relationships only with a structured dataset. The predictive results are based on low or complex interactions within features, including job role, level of satisfaction, and job duration.

Consequently, deep scaled MLP behave sustainable and achieved 96.68%, indicating deep learning strength for extracting representations without manual feature extraction. Such improved results indicated deep learning model can effectively learn hierarchical features and correlated underlying turnover and employee behavior.

The TabNet integrated attention mechanisms with feature selection using the deep learning architecture. It achieved 99.57% accuracy highest among all available TabNet algorithms. The highest performance explores attention based architecture suited for organized tabular data. The feature mask enables an interpretable TabNet. Accuracies of proposed models are shown in Table 5.5.

Table 5. 5: Accuracies of proposed models

Model	Accuracy	Precision	Recall	F1-Score
Ensemble decision tree	86.05%	Class 0: 87% Class 1: 68.75%	Class 0: 98% Class 1: 23.4%	Class 0: 92.2% Class 1: 34.9%
Tuned KNN	86.39%	Class 0: 86.83% Class 1: 76.92%	Class 0: 98.79% Class 1: 21.28%	Class 0: 92.42% Class 1: 33.33%
Deep scaled MLP	96.68%	Class 0: 87.74% Class 1: 45.45%	Class 0: 92.71% Class 1: 31.91%	Class 0: 90.16% Class 1: 37.50%
TabNet	99.57%	Class 0: 90.94% Class 1: 60%	Class 0: 93.52% Class 1: 51.06%	Class 0: 92.22% Class 1: 55.17%

Overall results confirmed that deep learning outperformed machine learning algorithms with the IBM HR dataset. Whereas machine learning models are less explainable and simple. The deep learning models with improved insight features are considered suitable with HR decision making.

6. Conclusion and Future Work

The TabNet achieved the best accuracy with 99.57% that is the best among all four models. The Deep Scaled Multi-Layer Perceptron obtained 96.68% accuracy, ranked second. Similarly, Tuned KNN performed well than Ensemble Decision Tree with 86.39% and 86.05% accuracies, respectively. The predictions and accuracies clearly confirmed that deep learning models, especially TabNet, effectively captured nonlinear dependencies as well as complicated among HR features for turnover impact. Similarly, TabNet interpretable behavior, such as feature attention enable it to deploy in real-life decision-making, especially in HR. While machine learning explored transparency and simplicity. Deep learning models provides better analytical performance and scalability constructing it a preferred choice for employee turnover prediction.

6.1. Limitations

The limitations of this work includes.

1. The results are specifically based on the IBM HR dataset that may vary with different with other datasets or organizational context
2. The categorical features may have imbalance and bias, caused affecting model generalization
3. The dataset is static, which may affect model behavior over turnover

6.2. Future works

Future research may aim to explore more scalable, robust, and interpretable prediction models for turnover. The future studies can be focused on: -

1. Implement longitudinal datasets with temporal models such as LSTM for tracking employee lifecycle, and may include multiple data sets.
2. Deploy models in practical life on the dashboard of HR to improve strategies and decision-making, and try role-based predictions to avoid biases.

References

- Ahmed, S.R., Ahmed, A.K. and Jwmaa, S.J. (2023) *Analyzing The Employee Turnover by Using Decision Tree Algorithm _ IEEE Conference Publication _ IEEE Xplore, IEEE Xplore.*
- Al-Suraihi, W.A. *et al.* (2021) ‘Employee Turnover: Causes, Importance and Retention Strategies’, *European Journal of Business and Management Research*, 6(3), pp. 1–10. Available at: <https://doi.org/10.24018/ejbmr.2021.6.3.893>.
- Alqahtani, H., Almagrabi, H. and Alharbi, A. (2024) ‘EMPLOYEE ATTRITION PREDICTION USING MACHINE LEARNING MODELS: A REVIEW PAPER’, 15(2). Available at: <https://doi.org/10.5121/ijaia.2024.1520223>.
- An, M. *et al.* (2022) ‘Factors Affecting Turnover Intention among New Graduate Nurses: Focusing on Job Stress and Sleep Disturbance’, *Healthcare (Switzerland)*, 10(6), pp. 4–12. Available at: <https://doi.org/10.3390/healthcare10061122>.
- Angermann, C. *et al.* (2024) ‘Three-Dimensional Bone-Image Synthesis with Generative Adversarial Networks’, *Journal of Imaging*, 10(12). Available at: <https://doi.org/10.3390/jimaging10120318>.
- Arqawi, S.M. *et al.* (2022) ‘Predicting Employee Attrition and Performance Using Deep Learning’, *Journal of Theoretical and Applied Information Technology*, 100(21), pp. 6526–6536.
- Atef, M., Elzanfaly, D.S. and Ouf, S. (2022) ‘Early Prediction of Employee Turnover Using Machine Learning Algorithms’, *International Journal of Electrical and Computer Engineering Systems*, 13(2), pp. 135–144. Available at: <https://doi.org/10.32985/IJECES.13.2.6>.
- BALAZS GONCZY (2022) ‘TABULAR EMPLOYEE RETENTION PREDICTION’, pp. 1–35.
- BALCIOĞLU, Y.S. and ARTAR, M. (2022) ‘Employee Turnover Prediction With Machine Learning: an Application With the K-Nearest Neighbors’, (February), pp. 106–111. Available at: <https://www.researchgate.net/publication/358575764>.
- Bustillo, J.C.M. (2025) ‘Optimization-based Techniques Prediction Model in Determining Employee Turnover’, *Procedia Computer Science*, 252, pp. 440–449. Available at: <https://doi.org/10.1016/j.procs.2025.01.003>.

- Chakraborty, R. *et al.* (2021) 'Study and Prediction Analysis of the Employee Turnover using Machine Learning Approaches', *2021 IEEE 4th International Conference on Computing, Power and Communication Technologies, GUCON 2021* [Preprint]. Available at: <https://doi.org/10.1109/GUCON50781.2021.9573759>.
- Chung, D.J., Kim, A. and Kim, Y. (2022) 'The contagion effect of collective voluntary turnover on firm performance and moderation of communication practices', *Human Resource Management Journal*, 32(1), pp. 19–39. Available at: <https://doi.org/10.1111/1748-8583.12346>.
- Durães, N.R. de A. (2024) 'Master Degree Program in Information Management Employee Turnover Predictive Model', pp. 1–80.
- Fallucchi, F. *et al.* (2020) 'Predicting employee attrition using machine learning techniques', *Computers*, 9(4), pp. 1–17. Available at: <https://doi.org/10.3390/computers9040086>.
- G. Klop (2021) 'Predicting Employee Turnover and Reducing Turnover Costs Using Machine Learning Techniques', *Industrial Engineering & Innovation Sciences*, p. 98. Available at: https://pure.tue.nl/ws/portalfiles/portal/183311084/Master_Thesis_Guido_Klop.pdf.
- Iparraguirre-Villanueva, O. *et al.* (2024) 'Employee attrition prediction using machine learning models', *Proceedings of the LACCEI international Multi-conference for Engineering, Education and Technology* [Preprint], (October). Available at: <https://doi.org/10.18687/LACCEI2024.1.1.498>.
- Karyati, S. *et al.* (2025) 'Exploring the Relationship between Generational Diversity , Employee Engagement , and Turnover Intention : Evidence from Manufacturing Workers', 03(02), pp. 86–95.
- Krishna, S. and Sidharth, S. (2022) 'HR Analytics: Employee Attrition Analysis using Random Forest', *International Journal of Performability Engineering*, 18(4), pp. 275–281. Available at: <https://doi.org/10.23940/ijpe.22.04.p5.275281>.
- Lazzari, M., Alvarez, J.M. and Ruggieri, S. (2022) 'Predicting and explaining employee turnover intention', *International Journal of Data Science and Analytics*, 14(3), pp. 279–292. Available at: <https://doi.org/10.1007/s41060-022-00329-w>.
- Maloku, F. (2024) 'Analyzing IBM HR Data: Employee Attrition and Performance Insights', *Journal of Engineering and Applied Sciences Technology*, 2024(August), pp. 1–10. Available at: [https://doi.org/10.47363/jeast/2024\(6\)268](https://doi.org/10.47363/jeast/2024(6)268).
- Munasinghe, L. and Sigman, K. (2004) 'A hobo syndrome? Mobility, wages, and job turnover', *Labour Economics*, 11(2), pp. 191–218. Available at: <https://doi.org/10.1016/j.labeco.2003.05.001>.
- Najafi-Zangeneh, S. *et al.* (2021) 'An improved machine learning-based employees attrition prediction framework with emphasis on feature selection', *Mathematics*, 9(11). Available at: <https://doi.org/10.3390/math9111226>.
- Park, J., Feng, Y. and Jeong, S.P. (2024) 'Developing an advanced prediction model for new employee turnover intention utilizing machine learning techniques', *Scientific Reports*, 14(1), pp. 1–14. Available at: <https://doi.org/10.1038/s41598-023-50593-4>.
- Prechsl, S. (2025) 'Reciprocity and job mobility : The effect of effort-reward imbalance in the employer-employee relationship on turnover intentions and actual job changes', *Social Science Research*, 127(October 2024), p. 103133. Available at: <https://doi.org/10.1016/j.ssresearch.2024.103133>.
- Qutub, A. *et al.* (2021) 'Prediction of Employee Attrition Using Machine Learning and Ensemble Methods', *International Journal of Machine Learning and Computing*, 11(2), pp. 110–114. Available at: <https://doi.org/10.18178/ijmlc.2021.11.2.1022>.
- Ramasamy, R.K., Muniandy, M. and Subramanian, P. (2025) 'A Predictive Framework for Sustainable Human Resource Management Using tNPS-Driven Machine Learning

- Models', pp. 1–28.
- Raza, A. *et al.* (2022) 'applied sciences Approaches'.
- Reddy, J.V. (2023) 'Big data-based framework for prediction of employee attrition by using deep data people analytics', *Journal of Informatics Education and Research*, 3(2), pp. 2547–2553. Available at: <https://doi.org/10.52783/jier.v3i2.432>.
- Seth, K. (2025) 'Predicting Telecom Customer Churn Utilising Machine Learning Models', *2025 World Skills Conference on Universal Data Analytics and Sciences (WorldSUAS)*, pp. 1–6. Available at: <https://doi.org/10.1109/WorldSUAS66815.2025.11198977>.
- Shiri, F.M. and Yamaguchi, S. (2025) 'A Deep Learning Model Based on Bidirectional Temporal Convolutional Network (Bi-TCN) for Predicting Employee Attrition'.
- Urme, U.N. (2023) 'The Impact of Talent Management Strategies on Employee Retention', *International Journal of Science and Business*, 28(1), pp. 127–146. Available at: <https://doi.org/10.58970/ijsb.2209>.
- Venkat, B.T.S., Chaitanya, J.V. and Teja, S.V.S. (2023) *EMPLOYEE CHURN PREDICTION*, *IEEE Xplore*.
- Yahia, N. Ben, Hlel, J. and Colomo-Palacios, R. (2021) 'From Big Data to Deep Data to Support People Analytics for Employee Attrition Prediction', *IEEE Access*, 9, pp. 60447–60458. Available at: <https://doi.org/10.1109/ACCESS.2021.3074559>.
- Živković, A. and Šebalj, D. (2024) 'Prediction of the Employee Turnover Intention Using Decision Trees Decision Trees and Employee', 2(Iceis), pp. 978–989. Available at: <https://doi.org/10.5220/0012538400003690>.