

Article

A Deep Learning Model Based on Bidirectional Temporal Convolutional Network (Bi-TCN) for Predicting Employee Attrition

Farhad Mortezapour Shiri ¹, Shingo Yamaguchi ^{2,*} and Mohd Anuaruddin Bin Ahmadon ³

¹ Faculty of Computer Science and Information Technology, University Putra Malaysia (UPM), Serdang 43400, Malaysia; gs63904@student.upm.edu.my

² Graduate School of Sciences and Technology for Innovation, Yamaguchi University, Ube 7558611, Japan

³ Department of Computer and Information Sciences, Universiti Teknologi PETRONAS, Seri Iskandar 32610, Malaysia; anuaruddin.ahmadon@utp.edu.my

* Correspondence: shingo@yamaguchi-u.ac.jp

Abstract: Employee attrition, which causes a significant loss for an organization, is the term used to describe the natural decline in the number of employees in an organization as a result of numerous unavoidable events. If a company can predict the likelihood of an employee leaving, it can take proactive steps to address the issue. In this study, we introduce a deep learning framework based on a Bidirectional Temporal Convolutional Network (Bi-TCN) to predict employee attrition. We conduct extensive experiments on two publicly available datasets, including IBM and Kaggle, comparing our model's performance against classical machine learning, deep learning models, and state-of-the-art approaches across multiple evaluation metrics. The proposed model yields promising results in predicting employee attrition, achieving accuracy rates of 89.65% on the IBM dataset and 97.83% on the Kaggle dataset. We also apply a fully connected GAN-based data augmentation technique and three oversampling methods to augment and balance the IBM dataset. The results show that our proposed model, combined with the GAN-based approach, improves accuracy to 92.17%. We also applied the SHAP method to identify the key features that most significantly influence employee attrition. These findings demonstrate the efficacy of our model, showcasing its potential for use in various industries and organizations.



Academic Editor: Stefan Fischer

Received: 5 February 2025

Revised: 4 March 2025

Accepted: 7 March 2025

Published: 10 March 2025

Citation: Mortezapour Shiri, F.; Yamaguchi, S.; Ahmadon, M.A.B. A Deep Learning Model Based on Bidirectional Temporal Convolutional Network (Bi-TCN) for Predicting Employee Attrition. *Appl. Sci.* **2025**, *15*, 2984. <https://doi.org/10.3390/app15062984>

Copyright: © 2025 by the authors. Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<https://creativecommons.org/licenses/by/4.0/>).

Keywords: employee attrition; employee turnover; deep learning model; bidirectional temporal convolutional network (Bi-TCN); data augmentation; generative adversarial network (GAN); explainable AI (XAI)

1. Introduction

In recent years, human resources (HR) management has garnered increasing attention due to its critical role in fostering competitive advantage and promoting organizational growth [1]. It is widely recognized that employees constitute an organization's most valuable asset, with the overall success of the company being closely tied to their knowledge, skills, and experience.

Employee attrition, or employee turnover, has become a focal point for human resources (HR) professionals, as it serves as a key indicator of organizational competitiveness. Attrition refers to the process through which employees leave a company, necessitating the recruitment of replacements. This phenomenon affects businesses across various industries and has been extensively studied in the recent literature [2]. Factors contributing to employee attrition include personal reasons, job dissatisfaction, inadequate compensation, and unfavorable work conditions, among others.

Employee attrition can be classified into two distinct categories: involuntary and voluntary. Involuntary attrition occurs when employees are terminated due to reasons such as poor performance or organizational restructuring. Conversely, voluntary attrition refers to situations where high-performing employees choose to leave the organization, often despite efforts by the company to retain them [3].

Employee attrition can pose a significant challenge to an organization's competitive advantage, leading to a variety of negative consequences. This issue has garnered considerable attention across industries, primarily due to its detrimental impact on organizational performance, operational efficiency, and the continuity of long-term growth initiatives. Attrition can incur substantial costs, as companies must devote extra resources for hiring, and training of new staff [4,5]. Companies must dedicate a lot of resources and time to training each employee according to the company's needs. When an employee leaves, not only does the organization lose valuable talent, but it also forfeits the considerable investments made in recruitment, screening, and training. Consequently, to fill vacancies, the company must once again dedicate significant resources to hiring, training, and integrating new staff members, perpetuating the cycle of cost and disruption [6]. Figure 1 illustrates the key aspects of human resources (HR) management and the impact of employee attrition on these areas.



Figure 1. The scope of human resources (HR) management and the impact of employee attrition on these domains.

Given these challenges, minimizing employee turnover becomes a central objective for organizations striving to maintain a stable and productive workforce. Companies can mitigate the negative effects of attrition on performance, morale, and overall business outcomes by fostering more engaging work environments and implementing effective organizational policies. By adopting targeted strategies to reduce attrition, businesses can not only lower the costs associated with recruitment and training but also enhance their long-term competitive edge and improve workforce stability [2].

In the contemporary digital landscape, business strategies are increasingly shaped by the integration of cutting-edge technologies like artificial intelligence (AI), machine learning (ML), and deep learning (DL). These technologies have evolved beyond their initial supporting roles, now playing a central role in contemporary commercial systems.

They change how strategic decisions are made in addition to increasing operational efficiency [7,8]. Today, nearly every industry stands to benefit from the adoption of these cutting-edge technologies. The capabilities for data collection, management, and analysis provide substantial advantages, boosting productivity and strengthening competitive positioning [9].

Machine learning algorithms offer a powerful tool for predicting employee attrition by analyzing factors such as job satisfaction, engagement levels, career progression, and other relevant variables. This predictive capability enables HR managers to implement targeted interventions designed to retain top talent. An essential component of enhancing decision-making processes is machine learning, which stands at the forefront of the data science field. These algorithms are designed to surpass human accuracy, continuously learning and adapting to new data inputs to make more informed predictions [10]. Machine learning in artificial intelligence (AI) is broadly categorized into two primary types: classical machine learning (CML) and deep learning (DL). Both approaches leverage historical data to enable machines to learn and make future predictions, with deep learning models often offering more complex, nuanced insights.

In this study, we introduce a deep learning framework for predicting employee attrition. The main contributions of our research are as follows:

- We propose an ensemble deep learning model that leverages a Bidirectional Temporal Convolutional Network (Bi-TCN) for employee attrition classification.
- To enhance model performance and address data imbalances, we incorporate a fully connected GAN-based data augmentation technique, which not only balances the dataset but also increases the volume of training data.
- We implement several baseline classical machine learning models, along with a few baseline deep learning models, and compare their results with the proposed model to better illustrate its effectiveness.
- We also apply the SHAP method, an explainable AI technique, to the proposed model to identify the key features that have the most significant impact on employee attrition.

The paper is structured as follows: Section 2 offers a thorough review of recent advancements in employee attrition prediction, emphasizing key methodologies and their contributions to the field. In Section 3, we explain the proposed deep learning methodology, detailing its architectural components and the rationale behind its design. Section 4 presents experimental results, including performance evaluations and comparative analyses of the proposed model. The work is finally concluded in Section 5, which offers a summary of the main conclusions and possible directions for further study.

2. Related Works

Employee attrition has garnered significant attention from researchers, management teams, and human resources professionals due to its potential to impact organizational stability. High attrition rates can lead to a number of detrimental effects, including increased recruitment and training costs, diminished team cohesion, loss of institutional knowledge, and disruptions to workflows [4]. Consequently, recent research has explored and applied a range of machine learning techniques to predict staff attrition. This section provides an in-depth review of the research on various employee attrition models, examining their effectiveness in predicting turnover and highlighting the efforts made to design robust classifiers for forecasting employee attrition.

Numerous studies have examined and evaluated various machine learning algorithms for predicting employee turnover [11–13]. Zhao et al. [14] evaluated ten supervised machine learning techniques across two datasets: IBM HR Analytics and a regional bank dataset. Along with statistical analysis, a number of data mining methods were used, including

cross-validation, parameter tuning, and data scaling. The results demonstrated that the best-performing models overall were tree-based classifiers, particularly XGB, GBT, RF, and DT. In another study [15], the performance of several machine learning models was evaluated on different feature subsets to predict employee attrition using the IBM dataset. Initially, five base models were trained and evaluated. Subsequently, three ensemble models were constructed by combining these base models in various ways. The results demonstrated the superior performance of the linear model in terms of AUC (area under the ROC curve), recall, and accuracy. A study by [9] utilized machine learning models to determine the characteristics that lead to employee turnover and, more crucially, to forecast the probability that a particular employee will leave the company. The methods applied in this research included Decision Tree (DT), Gaussian Naive Bayes, Support Vector Machine (SVM), K-Nearest Neighbors (KNNs), Logistic Regression (LR), Random Forest (RF), Linear Support Vector Machine (L-SVM), and the Naive Bayes classifier for multivariate Bernoulli models. The results indicated that the Gaussian Naive Bayes classifier achieved the highest F1-score (44.6%) on the IBM dataset, while the Linear SVC method delivered the best accuracy (87.9%).

Several studies have specifically focused on the use of Random Forest (RF) for forecasting employee turnover [16–18]. Gao et al. [19] introduced an innovative approach leveraging an enhanced Random Forest algorithm called weighted quadratic random forest (WQRF) to improve the predictive capacity for employee turnover. The proposed method is designed to address high-dimensional, imbalanced data. The approach begins by reducing the dimensions and ranking feature importance using Random Forest. Next, the Random Forest technique is applied to the selected features, with F-measure values computed as weights for each decision tree to construct the turnover prediction model. In another study, Jain et al. [20] presented a new tree-based ensemble method called eXplainable Reasonably Randomized Forest algorithm (XRRF), which balances model interpretability, accuracy, and explainability. They first introduced a graph-based feature learning (SGFL) technique, enhancing model accuracy by capturing feature co-dependencies while maintaining interpretability. This was followed by the proposal of a Reasonably Randomized Forest (RRF) method, which forms part of the XRRF ensemble. To further clarify the model's results, the decision-path feature extraction approach was also introduced. Additionally, a study by [21] applied an extensively tuned principal component analysis (PCA) technique in conjunction with Random Forest classifiers to predict employee turnover. In this work, PCA was used to identify relevant features, while the Random Forest classifier was employed for prediction. The framework's efficacy in predicting employee departure based on job attitudes and other internal and external factors was demonstrated through a comparison of various performance metrics.

Several researchers have utilized Logistic Regression (LR) for predicting employee turnover [5,22]. Najafi-Zangeneh et al. [23] proposed a three-stage approach for attrition prediction, encompassing pre-processing, processing, and post-processing phases. In the pre-processing phase, the authors introduced the m-max-out method for feature selection. Logistic Regression (LR) was then employed as the classification technique. The model's validity was assessed by examining the changes in its parameters across multiple bootstrap datasets, verifying the robustness of the LR model for attrition prediction. These pre-processing and post-processing stages contribute to the development of more accurate and reliable models for employee turnover forecasting.

Several studies have demonstrated that boosting algorithms tend to outperform other machine learning techniques when it comes to forecasting employee turnover [24,25]. Atique et al. [26] employed an enhanced feature engineering approach in conjunction with the boosting algorithm CatBoost to predict and analyze employee turnover using

the IBM dataset. According to experimental results, the proposed method outperformed existing models. Specifically, it achieved an accuracy of 89.45% and an F1-score of 88.0% for employee attrition prediction on the IBM dataset. In another study, Jain et al. [27] proposed an approach based on multi-attribute decision-making (MADM) and machine learning algorithms, referred to as “Employee Classification and Prediction for Retention” (ECPR). Using a two-stage MADM technique, they developed an accomplishment-based employee importance model (AEIM) to divide staff into different groups. To assign relative weights to employee accomplishments, they introduced an enhanced version of the entropy weight method (IEWM). Employee performance significance within each class was then measured using the order preference by similarity to ideal solution (TOPSIS) approach. Subsequently, employee turnover was predicted for each class using the CatBoost algorithm. Finally, a retention strategy was proposed based on the feature ratings and the predictive findings.

Several scholars have focused on utilizing deep learning and neural network models to predict employee turnover [28,29]. Al-Darraji et al. [30] employed a Deep Neural Network (DNN) combined with several pre-processing techniques to improve employee attrition prediction. Their model obtained an accuracy of 89.11% using 10-fold cross-validation on the original IBM dataset. To enhance the realism of the results, they also created a balanced version of the dataset, which led to an accuracy increase to 94%. In a separate study [31], deep learning approaches were explored to enhance the accuracy of employee turnover prediction. The authors proposed a multi-layered neural network architecture that integrates data from diverse sources, including demographics, employee engagement metrics, and historical turnover data. By combining feedforward and recurrent networks, the model effectively captures complex relationships and temporal dependencies in the data. Furthermore, advanced feature engineering techniques were employed to transform raw data into valuable inputs, significantly boosting the model’s predictive performance. The experimental results indicated that the proposed approach outperformed traditional machine learning methods in terms of accuracy and reliability. Furthermore, Mohamed Ahmed [32] developed a novel data mining model that incorporates Information Gain and Chi-Square as feature selection techniques. These methods were used to identify the four most significant features in the dataset: overtime, job level, salary, and years in the organization. Several classification algorithms, including Decision Tree, SVM, Random Forest, Neural Network, and Naïve Bayes, were applied to construct the model. Based on the implementation results on the IBM dataset, the Neural Network approach emerged as the most successful, obtaining an accuracy of 84%.

Additionally, the authors of [33] proposed an event-centered turnover prediction method called CoxRF, which integrates ensemble learning with the statistical insights of survival analysis. They completely used restricted data by introducing the ideas of “event-person” and “time-event” to generate survival statistics. Their findings highlight several key insights: (i) Gender is a significant factor influencing staff attrition behavior, with the female staff exhibiting a higher attrition rate than their male counterparts; (ii) external factors, like the growth of Gross domestic product (GDP), have a notable impact on employee turnover, a consideration that has been largely overlooked in most studies; (iii) Different industries have different trends in staff attrition; the IT industry has a far greater rate than the government sector; (iv) Strongly educated staff typically leave their jobs more frequently than those with less education, especially after three to five years.

Jin et al. [34] proposed a turnover prediction technique called RFRSE, which develops a hybrid model that integrates machine learning with survival analysis. This approach combines ensemble learning for predicting turnover behavior with survival analysis to handle censored data. In order to create survival statistics using restricted information, the authors also implement methods to address employees with multiple turnover records.

Using a real dataset taken from one of China's biggest professional social media platforms, they compare the performance of RFRSF with several baseline techniques. The findings demonstrate that the survival analysis component greatly enhances the effectiveness of employee attrition prediction.

The Random Forest (RF) and K-Nearest Neighbors (KNNs) methods served as the foundation for the models put out by [35]. According to the experimental results, the KNN-based method outperformed the RF-based method, achieving an accuracy of 84.0% compared to 80.0% on the IBM dataset.

Al Akasheh et al. [36] introduced a unique method that utilizes Graph Convolutional Networks (GCNs) to convert typical tabular employee data into a knowledge graph structure in order to extract more subtle information. The method integrates both graph-derived information and the original IBM dataset to predict employee turnover. Several machine learning models were employed to evaluate classification performance across various criteria, with the results indicating that the Linear Support Vector Machine (L-SVM) emerged as the most effective model.

A hybrid model combining an Autoencoder, a Genetic Algorithm, and K-Nearest Neighbor was proposed by [37] for forecasting employee turnover, named GA-DeepAutoencoder-KNN. The approach enhances prediction accuracy by integrating the KNN model, an Autoencoder, and a Genetic Algorithm. The model was empirically assessed and contrasted with standard KNN and DeepAutoencoder-KNN methods. The results demonstrated that, using the IBM dataset, the GA-DeepAutoencoder-KNN method obtained an accuracy of 90.95%, outperforming the DeepAutoencoder-KNN model (86.48%) and the KNN model (88.37%).

A few studies also focus on explainable AI (XAI) methods to identify the most influential features and provide valuable insights for HR decision-making. Díaz et al. [38] explored how explainable AI (XAI) can be used to detect possible staff attrition and develop data-driven solutions to deal with this challenging issue. Initially, they concentrated on using machine learning models for predicting employee attrition. Then, in order to improve the transparency and interpretability of AI models, they used explainable AI (XAI) techniques like SHAP (SHapley Additive exPlanation) [39], and LIME (Local Interpretable Model-agnostic Explanations) [40]. The authors of [41] used the Explainable Graph Neural Network (GNN), and Graph Attention Network (GAT) to predict employee leave and pinpoint key variables influencing their choices. They exploited the GNN model's capacity to identify the deep-rooted structure of employee data, where linkages between coworkers can hold significant insights. By using explainable AI methods, the model highlights the most significant factors influencing employee turnover and produces interpretable predictions. In another study, the authors of [42] introduce a novel approach to predict employee turnover by combining clustering techniques with Artificial Neural Networks (ANNs). To obtain the best ANN models, they concentrate on hyperparameter tuning with different input parameters. The study's data segmentation helps identify important turnover predictors, which enables the implementation of focused interventions to increase the efficacy and efficiency of retention strategies. Conditional Generative Adversarial Networks (CTGANs) are used to augment data on clusters that have unbalanced data. After that, these augmented clusters are subjected to the optimized ANN models, which significantly enhances model performance. The SHAP method was utilized to assess each feature's significance across various clusters in a predictive model. Varkiani et al. [43] evaluated four machine learning models for employee attrition prediction using a real dataset obtained from an Italian financial company. They discovered that the model with the greatest performance was Random Forest. They also used the ROSE [44] technique in combination with Random Forest to address the class imbalance problem because of the

very unbalanced dataset. Moreover, a SHAP (SHapley Additive exPlanation) technique was used to identify feature contributions and assess their direction.

Despite considerable efforts in predicting employee attrition, this research domain continues to face significant challenges that warrant further investigation. Many existing studies predominantly rely on classical machine learning models or relatively basic deep learning architectures, which may fall short in effectively capturing the complexities of employee attrition data. To achieve more accurate predictions, the development and application of advanced, robust models are essential. Additionally, the pre-processing stage, particularly data augmentation, is often overlooked, despite its crucial role in improving model performance and reliability. Furthermore, only a few articles utilized explainable AI (XAI) methods to identify the most influential features and offer valuable insights into HR decision-making.

3. Methodology and Proposed Model

Recent developments in deep learning (DL) have demonstrated impressive promise for large-scale dataset analysis to uncover complex patterns, offering significant improvements in predictive accuracy for employee turnover. DL is the process of learning hierarchical data representations using a neural network structure that includes several hidden layers that add to the network's depth. In DL algorithms, data flow through these layers sequentially, with each layer gradually extracting increasingly complex characteristics and forwarding essential data to the next. Low-level characteristics are captured by the first layers, which are subsequently integrated and improved by later layers to produce a thorough and intricate depiction of the data. Deep learning methods, renowned for their ability to automatically recognize and learn features from training data, have demonstrated superior performance across various classification tasks. Thanks to its inherent feature extraction capabilities, DL has become a frontrunner in tackling complex challenges, demonstrating a significant edge in addressing sophisticated problems [45].

In this study, we propose a novel ensemble deep learning framework designed to predict employee attrition with high accuracy. At the heart of this framework lies the Bidirectional Temporal Convolutional Network (Bi-TCN), which excels in capturing complex correlations within the data, significantly improving classification performance. Figure 2 illustrates a comprehensive overview of the proposed methodology and model architecture.

The ensemble consists of two Bi-TCN layers, each configured with distinct numbers of filters and kernel sizes to enhance feature extraction. In particular, 32 filters with a kernel size of 3 are used in the first Bi-TCN layer, while 64 filters with a kernel size of 5 are used in the second. This strategic utilization of varying kernel sizes and filter counts allows the model to capture a wider range of local dependencies, enabling a more comprehensive and detailed analysis of the data. Temporal Convolutional Networks (TCNs) can learn hierarchical feature representations because they use many convolutional layers with progressively larger receptive fields. Therefore, this method is very advantageous for non-sequential datasets with patterns across many feature dimensions. TCNs are also more successful on small or irregular datasets because they use weight sharing and dilation techniques, which decrease the number of trainable parameters and improve generalization.

We also utilize a batch normalization layer after the fully connected layer, two dropout layers with a rate of 0.5 after each Bi-TCN layer, and a weight decay regularization (L2) technique in order to prevent overfitting in the proposed deep learning model. Weight decay encourages the weights to be modest in size, which enhances generalization and lowers the chance of overfitting [46].

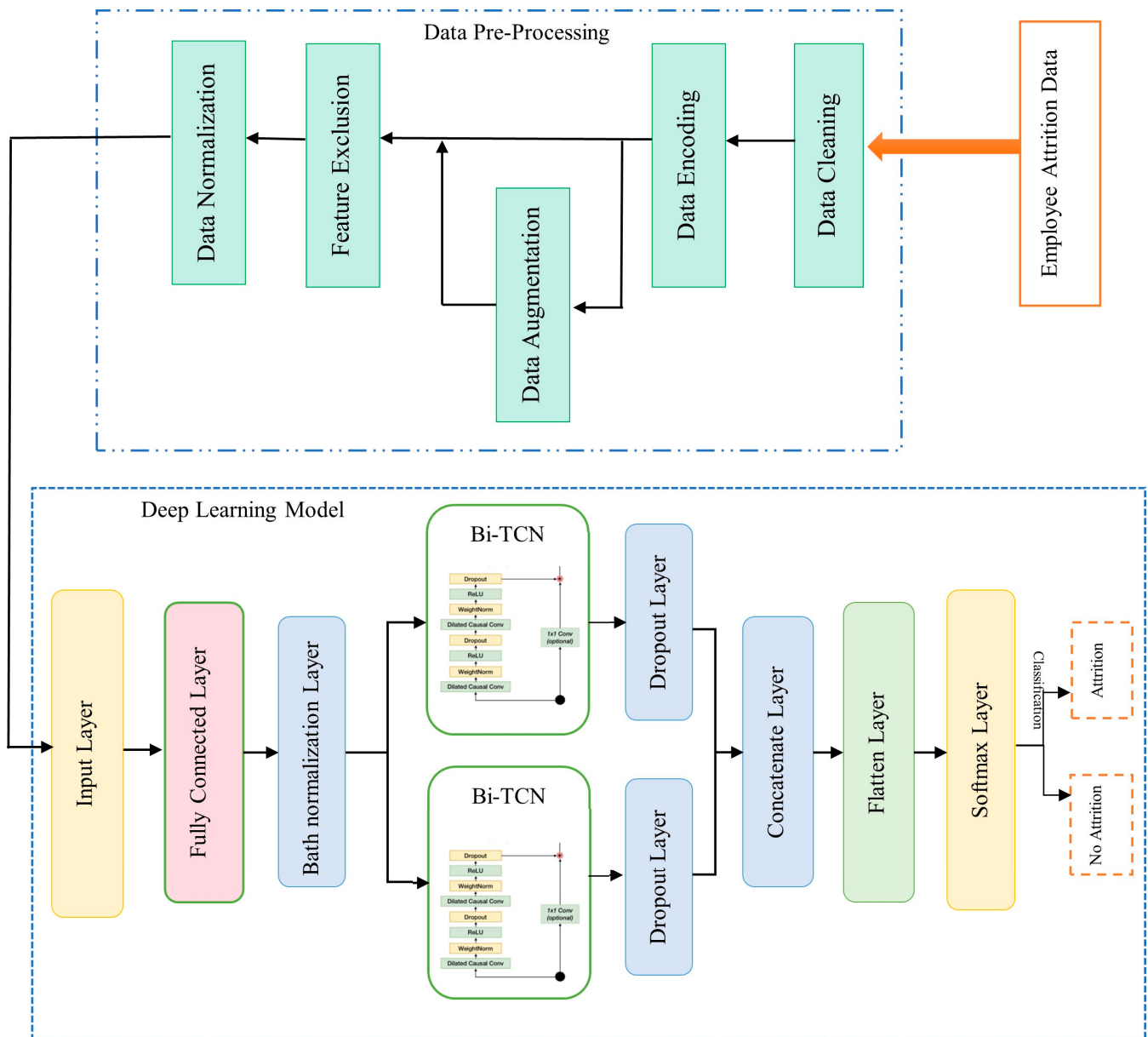


Figure 2. Proposed deep learning model for predicting employee attrition.

Temporal Convolutional Network (TCN): One notable advancement in deep learning architectures is the Temporal Convolutional Network (TCN) [47], which excels at analyzing temporal patterns while retaining the robust feature extraction capabilities of Convolutional Neural Networks (CNNs). TCN is characterized by two key features: (1) By using causal convolutions, it guarantees that the output at any given time step is entirely dependent on the inputs that are now and previously present, with no effect from inputs that may be added in the future. (2) Like Recurrent Neural Networks (RNNs), TCN can handle sequences of any length and generate outputs that match the input sequence in length. The typical TCN architecture comprises three main components: residual connections, dilated convolutions, and causal convolutions.

Causal convolutions are one-dimensional convolutional layers designed to use only data from time t and prior to that to compute the output at time t . Dilated convolutions expand the receptive field efficiently by skipping input elements at regular intervals, enabling the network to capture long-range dependencies without a significant increase in computational cost [48,49]. To create a more expressive model, multiple layers with

relatively small filter sizes are often stacked. Nevertheless, problems like vanishing or expanding gradients during training may arise if these dilated and causal convolutional layers are stacked to improve network depth. To overcome these challenges, TCN employs residual connections, which create direct pathways for data to bypass certain layers. These connections improve training stability and efficiency by allowing the network to learn residual functions that adjust the identity mapping instead of developing entirely new transformations [50]. A TCN model's schematic architecture is seen in Figure 3.

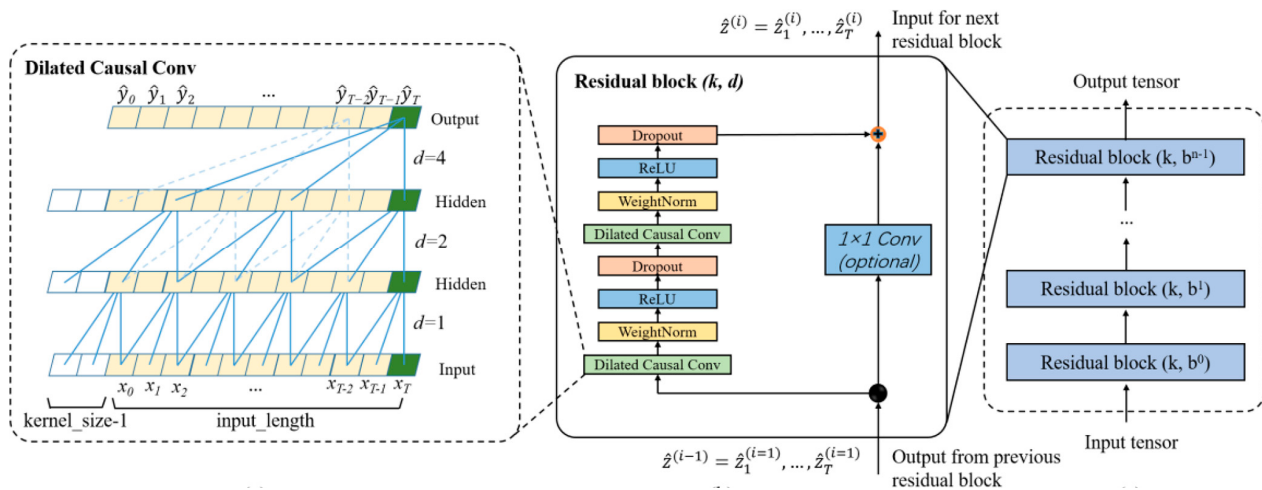


Figure 3. Schematic diagram of the TCN model [51].

Unlike standard TCN models, which process data solely in a forward direction, Bidirectional TCN (Bi-TCN) enhances performance by processing input data in both forward and backward orientations, enabling a more comprehensive analysis of temporal patterns.

3.1. Dataset

In this study, we employed the IBM HR Analytics dataset [52] and Kaggle Employee Churn Prediction dataset [53] to evaluate the effectiveness of our proposed deep learning method. The first dataset, developed by IBM Data Scientists, is specifically designed to identify the critical factors influencing employee turnover.

It contains 1470 records with 34 features, including attributes such as “age”, “gender”, “daily rate”, “job satisfaction”, and others, alongside a target column labeled “Attrition”. The “Attrition” column captures whether an employee has chosen to leave the company (“Yes”) or remain (“No”). Figure 4 provides a detailed visualization of the dataset’s features and their correlations.

To gain a clearer understanding of the most significant feature correlations, we filtered and highlighted correlations greater than 0.4 in an additional heatmap presented in Figure 5. For instance, the figure reveals strong correlations between ‘MonthlyIncome’ and both ‘JobLevel’ and ‘TotalWorkingYears’, while ‘PerformanceRating’ is highly correlated with ‘PercentSalaryHike’.

The second dataset is a big dataset offered by Kaggle that comprises 14,249 samples and 10 features. Details of the attributes are presented in Table 1.

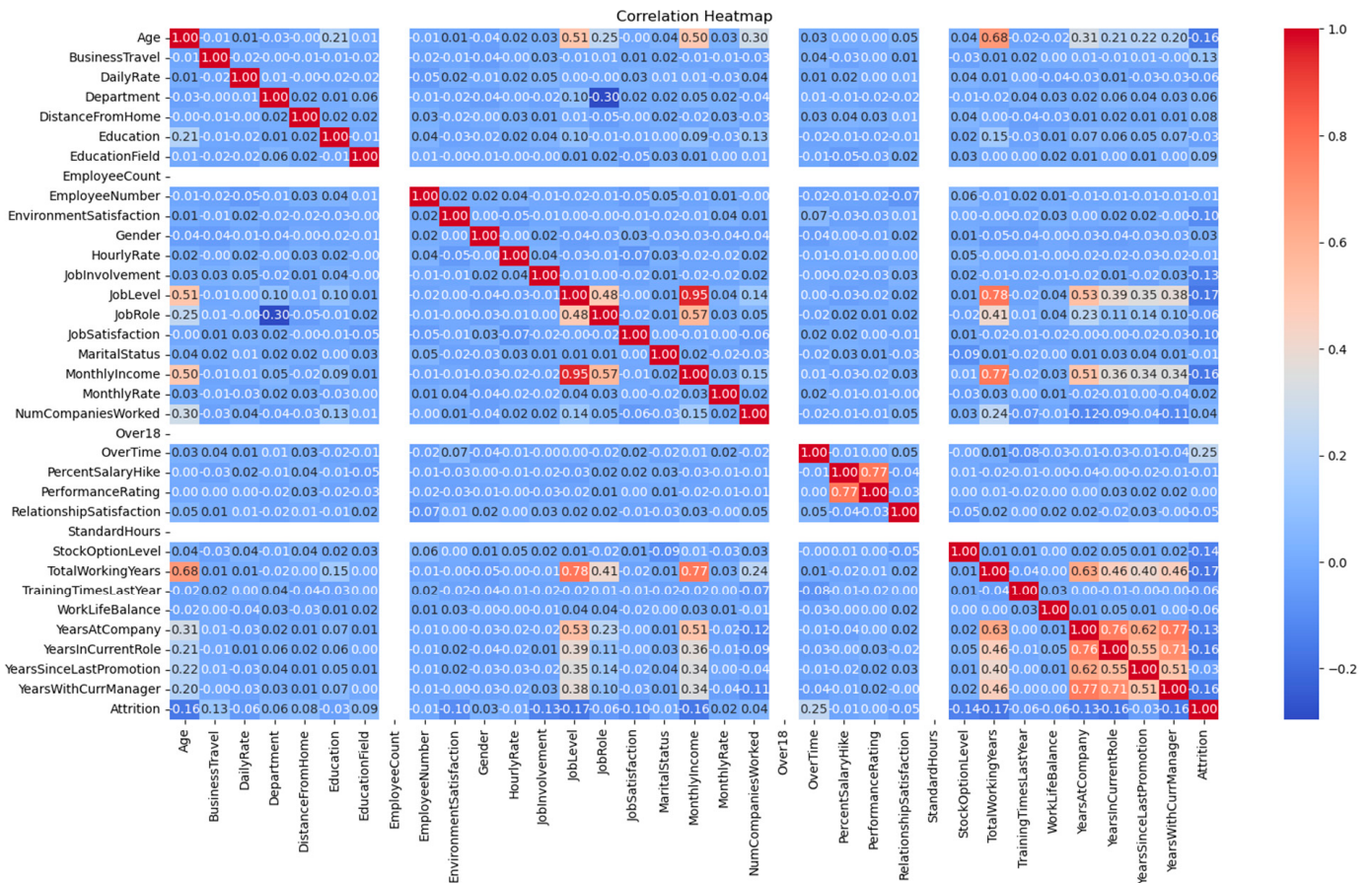


Figure 4. The correlation heatmap between the IBM dataset features.

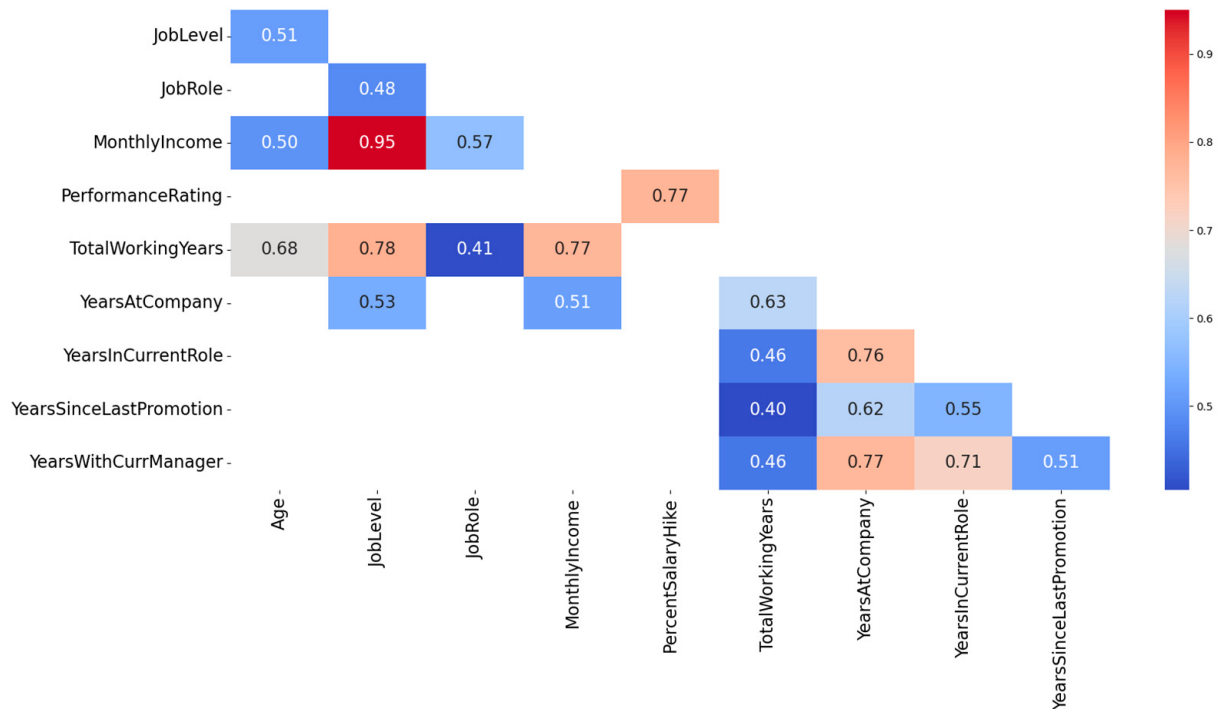


Figure 5. The correlation greater than 0.4 between the IBM dataset features.

Table 1. Attributes list of the Kaggle dataset.

Attribute	Data Type	Description
Avg-Monthly-Hrs	integer	The average number of hours worked per month.
Department	Categorical	The department of employee.
Filed-Complaint	integer	If the employee filed a formal complaint in the last 3 years.
Last-Evaluation	Float	Score for the most recent evaluation of the employee, which ranges from 0 to 1.
N-Projects	integer	The number of projects that an employee works on.
Recently-Promoted	integer	If a worker has received a promotion in the last three years.
Salary-Level	Categorical	The employee's salary level, categorized as low, medium, or high.
Satisfaction	Float	Score for employee satisfaction with the company, which ranges from 0 to 1.
Tenure	integer	Number of years in the company.
Status (target)	Categorical	Current employment status, (Employed/left).

3.2. Data Pre-Processing

Pre-processing is an essential stage in the data pipeline, involving various techniques and operations to transform raw data into a format suitable for analysis or machine learning model use. This subsection outlines the pre-processing steps, including data cleaning, data encoding, data augmentation, feature engineering, and data normalization, that will be applied to enhance the accuracy of employee attrition prediction.

3.2.1. Data Cleaning

The Kaggle dataset contains four attributes with missing values: Last-Evaluation (1532 missing values), Satisfaction (180 missing values), Department (706 missing values), and Tenure (180 missing values). We removed the instances with missing values for any of these attributes. Notably, the missing values for Satisfaction, Department, and Tenure occurred in records where Last-Evaluation was also missing. As a result, 1532 instances were removed, leaving 12,717 instances in the dataset.

3.2.2. Data Encoding

Since most machine learning algorithms cannot directly handle categorical features, these features must be transformed into numerical representations to be used in an ML model. Several categorical features, including "Business Travel", "Department", "Education Field", "Gender", "Job Role", "Marital Status", "Overtime", and "Attrition", are included in the original IBM dataset and must be transformed into numerical values. Table 2 shows the transformation of categorical feature values into numerical values.

Table 2. Transformation of categorical feature values into numerical values.

Categorical Feature	Description of Transformed into Numerical Values
Business Travel	0 = Non travel, 1 = Travel Rarely, 2 = Travel Frequently
Department	0 = Human Resources, 1 = Research and Development, 2 = Sales
Education field	0 = other, 1 = Life Sciences, 2 = Medical, 3 = Marketing, 4 = Life Sciences, 5 = Human Resources
Gender	0 = Female, 1 = Male
Job Role	0 = Sales Executive, 1 = Research Scientist, 2 = Laboratory Technician, 3 = Manufacturing Director, 4 = Healthcare Representative, 5 = human resources, 6 = Sales Representative, 7 = Research Director, 8 = manager
Marital Status	0 = Divorced, 2 = Married, 1 = Single
Over Time	0 = No, 1 = Yes

In the Kaggle dataset, there are three categorical features: “Department”, “Salary Level”, and “Status”. These features need to be converted into numerical values. For example, for the “Status” attribute, “employed” is represented as 0, while “left” is represented as 1.

3.2.3. Data Augmentation

Large amounts of data are essential for deep learning models to perform effectively, as increased data typically enhance model performance [54]. However, in certain cases, such as with the IBM dataset, the available data may be limited. With only 1470 entries, this dataset presents challenges for efficiently training deep learning models. Additionally, data imbalance poses another issue in the IBM dataset.

In binary classification tasks, class imbalance arises when the majority of the dataset’s examples fall into one class while the remainder fall into the other class [55]. In the IBM dataset, 1233 employees are in the “no” attrition group, while only 237 entries are in the “yes” attrition category, creating an unbalanced distribution between the two attrition groups.

One effective approach to enhancing the generalizability of trained models is data augmentation, which is especially valuable when working with small and imbalanced datasets, common challenges in practical applications. By adding synthetic instances to the dataset, classification accuracy can often be improved [56]. There are various data augmentation and data balancing techniques, each tailored to the specific features of the dataset. One effective method is the use of generative data augmentation techniques, such as Autoencoders (AEs) and Generative Adversarial Networks (GANs). Although Autoencoders have been explored for generating synthetic data, their adoption has been limited due to the lower quality of the data they produce. In contrast, GANs are capable of generating highly realistic synthetic instances, making them a more promising approach [57,58].

Generative Adversarial Networks (GANs): The two main parts of a GAN are a discriminative model and a generating model. While the generative model seeks to generate data that closely match actual data, the discriminative model seeks to distinguish between genuine and synthetic data [59,60]. Essentially, the GAN framework operates as a two-player adversarial game, where the discriminator (D) and the generator (G) compete with one another. The discriminator calculates the generator’s updating gradients using an adaptive objective [61]. Both components of the GAN may incorporate multiple deep learning layers, including fully connected layers, convolutional layers, and others.

We utilize a fully connected GAN-based data augmentation technique to generate a balanced synthetic dataset, significantly increasing the number of records compared to the original IBM dataset. Table 3 outlines the architectural details of the generator and discriminator in the GAN used in this study.

Table 3. The details of generator and discriminator layers of the GAN model.

Generator			Discriminator		
Layer	Configuration Details	Output	Layer	Configuration Details	Output
Input	Latent dim = 16	None, 16	Input	Input dim = 35	None, 35
Dense	Unit = 64	None, 64	Dense	Unit = 128	None, 128
LeakyReLU	Alpha = 0.2	None, 64	LeakyReLU	Alpha = 0.2	None, 128
Dense	Unit = 128	None, 128	Dense	Unit = 64	None, 64
LeakyReLU	Alpha = 0.2	None, 128	LeakyReLU	Alpha = 0.2	None, 64
Dense	Unit = 35, Activation = tanh	None, 35	Dense	Unit = 1, Activation = sigmoid	None, 1

Since the discriminator often outperforms the generator, this can prevent the training process from converging. To address this, it is important to set appropriate model parameters that do not overwhelm the generator. One key consideration when designing the discriminator is ensuring that its trainable parameters are roughly equal to those of the generator, thus maintaining a balance between both models [62].

We also applied three oversampling techniques, including the Random Oversampling, Synthetic Minority Oversampling Technique (SMOTE) [63], and Adaptive Synthetic Sampling (ADASYN) [64], to balance and augment the IBM data. Our goal is to compare these methods with each other and with the GAN-based method in terms of their effectiveness on the proposed model's performance for employee attrition prediction.

Synthetic Minority Oversampling Technique (SMOTE): SMOTE is an oversampling method that generates synthetic data points by using the K-Nearest Neighbors (KNNs) algorithm. In SMOTE, for each underrepresented instance, a certain number of nearest neighbors are found. Then, to create synthetic data points, a subset of minority class instances is chosen at random. Lastly, along the line segments that link the chosen minority examples to their nearest neighbors, new artificial observations are made [65].

Adaptive Synthetic Sampling (ADASYN): ADASYN is an oversampling algorithm that generates synthetic samples by applying a weighted distribution to different minority class instances. The way that ADASYN creates synthetic samples on the line segments between two minority class data points is similar to that of SMOTE. However, ADASYN automatically calculates how many synthetic samples to create for every instance of the minority class based on a density distribution. Therefore, a balanced representation of the data distribution is provided by the expanded dataset [66].

3.2.4. Feature Exclusion

The process of excluding specific features (variables) from a dataset while developing or analyzing a model is known as feature exclusion. By streamlining the model, concentrating on the most pertinent data, and minimizing overfitting, this method is frequently employed to enhance model performance. A quick review of the IBM dataset reveals that several features, such as "Employee Count", "Over18", and "Standard Hours", have identical values for all employees and have therefore been excluded. Furthermore, the "Employee Number" feature, as its values do not contribute to the objective of the analysis, has also been omitted.

3.2.5. Data Normalization

Real-world datasets often exhibit variations in range, units, and magnitude, which can lead to suboptimal classification outcomes. Features with wider ranges may dominate the model's learning process, overshadowing other important features [67]. For example, the "Daily Rate" feature in the IBM dataset spans from 234 to 2877, and this significant difference in values could impair model performance. To address this, it is essential to rescale the feature values to fall within a consistent range. A commonly used technique for rescaling is normalization. In this study, feature values have been rescaled to a range between 0 and 1, as shown in Equation (1). This normalization serves as a practical pre-processing step that enhances the model's performance, reduces bias, and improves model interpretability.

$$\hat{x} = \frac{x - x_{min}}{x_{max} - x_{min}} \quad (1)$$

where the values of the specified feature are denoted by x_{min} for the minimum and x_{max} for the maximum, respectively.

4. Experiments

For the experiments, we implemented the proposed deep learning model alongside several baseline classical machine learning models, including Decision Tree (DT), K-Nearest Neighbors (KNNs), Random Forest (RF), Logistic Regression (LR), Adaptive Boosting (AdaBoost), Gradient Boosting (GB), Extreme Gradient Boosting (XGBoost), and Categorical Boosting (CatBoost). Additionally, we employed deep learning models including Convolutional Neural Network (CNN), Multi-Layer Perceptron (MLP), Long Short-Term Memory (LSTM), Gated Recurrent Unit (GRU), Bidirectional LSTM (Bi-LSTM), and Bidirectional GRU (Bi-GRU). This comprehensive approach facilitates an efficient comparison and provides a clearer demonstration of the proposed model's performance. To familiarize readers with the structures and capabilities of these baseline models, we also provide concise explanations of some recent models.

4.1. Machine Learning Models

Random Forest (RF): The RF model [68] is an ensemble classifier that uses randomly chosen subsets of training data and features to build multiple decision trees. The classifier combines predictions from these individual decision trees to provide accurate and reliable classifications. Furthermore, RF can identify and rank the features that contribute most significantly to distinguishing between target classes, making it a powerful tool for feature importance analysis [69].

Logistic Regression (LR): This model is a straightforward parametric statistical method that constructs a model by identifying classification parameters capable of distinguishing between groups and formulating associated classification rules through numerical optimization. One of LR's key advantages is its ability to operate without making assumptions about class distributions in the feature space, enhancing its versatility in various classification tasks [2].

Adaptive Boosting (AdaBoost): AdaBoost [70] is an ensemble learning method that iteratively improves classification performance by focusing on the errors of weak classifiers. Unlike Random Forest, which uses parallel ensembling, AdaBoost employs "sequential ensembling", where each classifier is trained to correct the errors made by its predecessor. When it comes to improving decision tree performance for binary classification jobs, AdaBoost is especially useful. Nevertheless, its overall efficacy may be impacted by its sensitivity to outliers and noisy data [71].

Gradient Boosting (GB): Gradient Boosting (GB) is an ensemble model that builds a final predictive model by sequentially merging several independent models, usually decision trees. The method optimizes weights by leveraging gradients to minimize a specified loss function. By iteratively adding weaker models to correct the errors of the ensemble, Gradient Boosting creates a more robust predictor. This approach often outperforms single strong models in data-driven tasks, as it combines the strengths of weaker models into a cohesive and powerful ensemble estimator [72,73].

Extreme Gradient Boosting (XGBoost): XGBoost is a refined version of gradient boosting that incorporates more precise approximations to optimize the model. It employs the loss function's second-order gradients to decrease errors and includes sophisticated regularization techniques to reduce overfitting. These features enhance model performance and generalization. XGBoost is a preferred option for difficult data-driven tasks because of its quick learning capabilities and exceptional efficacy in managing huge amounts of data [71].

Categorical Boosting (CatBoost): CatBoost [74] is a gradient boosting technique designed specifically for handling categorical data effectively. Unlike other popular gradient boosting methods like XGBoost, CatBoost employs ordered boosting, ensuring that each

model in the ensemble is trained exclusively on historical data. This approach enhances generalization, and the likelihood of overfitting is decreased, making CatBoost a highly efficient and accurate option for datasets with significant categorical features.

4.2. Deep Learning Models

Multi-Layer Perceptron (MLP): One kind of feedforward Artificial Neural Network (ANN) that forms the basis of deep learning or the Deep Neural Network (DNN) is the MLP. An input layer, one or more hidden layers, and an output layer are the three primary components of MLP. Every neuron in a layer is connected to every other neuron in the layers that surround it, indicating that the MLP network is fully connected. The input layer receives and normalizes the features from the input data. The hidden layers, whose number might vary, use representations they have learned to process the input data. Finally, the output layer generates predictions or decisions based on the extracted data [48,75].

Convolutional Neural Networks (CNNs): One of the most effective models of deep learning is CNN. Leveraging a convolutional architecture, CNNs act as feedforward neural networks that automatically extract features from incoming data [76]. A classifier and a feature extractor are combined in CNN's two-stage architecture to allow automatic feature extraction and end-to-end training with minimal pre-processing. In contrast to conventional techniques, CNNs do not require manual feature engineering since they learn and recognize features straight from the data [48,77].

Long Short-Term Memory (LSTM): LSTM [78] is an advanced variant of Recurrent Neural Networks (RNNs) particularly designed to address the common issue of long-term dependencies. LSTM networks are highly effective in retaining information over long sequences and solving the vanishing gradient problem. An LSTM processes the current input and the output from the previous phase at each time step, resulting in an output that is transferred to the subsequent time step. The last hidden layer at the last time step is often used for classification. Three gates make up an LSTM: input, forget, and output gates. These gates are essential for efficiently handling reading and writing operations inside the LSTM architecture by controlling the flow of data into and out of the memory unit [79].

Bidirectional LSTM (Bi-LSTM): Bi-LSTM is an enhancement of the LSTM that captures both past and future context in sequence modeling tasks, addressing the limitations of LSTM designs. Unlike conventional LSTM, which processes input data in a single forward orientation, Bi-LSTM processes data in both forward and reverse orientations, allowing it to leverage additional contextual information from the entire sequence [79].

Gated Recurrent Unit (GRU): Another recurrent neural network (RNN) version that tackles the short-term memory problem in sequence modeling is the GRU [80], which has a simpler architecture than LSTM. With only two gates instead of the three LSTM gates and no cell state, it has a simpler design that allows for quicker learning because of its lower computational complexity. The GRU can efficiently handle data from previous time steps and record long-term situations in arrangements because of its design, which consists of a reset gate, an update gate, and the current memory value [81].

Bidirectional GRU (Bi-GRU): The Bi-GRU is an extension of the GRU design that incorporates both past as well as future instances in sequential modeling problems, hence overcoming some of the normal GRU's constraints. Unlike the GRU, which only examines input sequences in a forward orientation, the Bi-GRU can operate in both forward and reverse [81].

Transformer: The transformer architecture [82], widely recognized for its effectiveness in modeling sequential data, consists of multiple identical layers. A position-wise completely linked feedforward network and a multi-head self-attention mechanism are

integrated into each layer. Layer normalization and residual connections are added to these elements to stabilize training and promote gradient flow.

4.3. Dataset Splitting

The k-fold cross-validation method is employed to divide the dataset into subsets for training and testing during the experimental phase. k-fold cross-validation is a popular statistical technique that splits the whole dataset into k folds, which are equal-sized (or nearly equal) subgroups.

This approach enhances the robustness of model evaluation by minimizing reliance on a single train–test split. For each iteration, one fold serves as the test set, while the remaining k-1 folds are used for training. Each fold is used as the test set once during the k repetitions of the procedure. After completing all k iterations, the performance metrics are averaged across all folds to provide a comprehensive assessment [83,84]. Figure 6 illustrates the k-fold cross-validation process. By mitigating the influence of specific data splits, this method ensures a more reliable evaluation of the model’s performance on unseen data. The present investigation employs 5-fold cross-validation to produce accurate results and enable a thorough examination and comparison of experimental results.

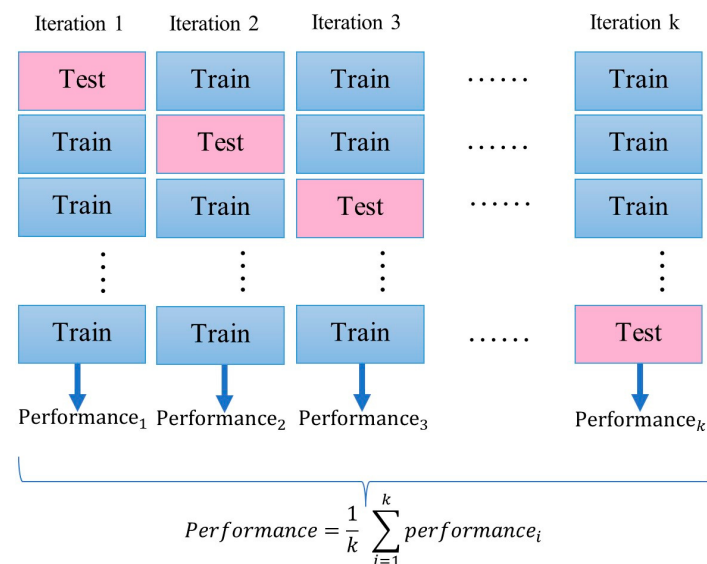


Figure 6. The k-fold cross-validation process for splitting data.

4.4. Evaluation Metrics

We utilized several key evaluation metrics, including accuracy, recall, precision, F1-score, and AUC (Area Under the ROC Curve), to thoroughly assess the performance of our proposed deep learning model and compare it with baseline machine learning and deep learning models.

- **Accuracy** provides a general assessment of the model’s performance by calculating the percentage of properly predicted cases relative to all instances, as illustrated in Equation (2).

$$Accuracy = \frac{Tp + Tn}{Tp + Tn + Fp + Fn} \quad (2)$$

- **Precision** evaluates how well the model prevents false positives by figuring out what proportion of all positive predictions are true positives, as depicted in Equation (3).

$$Precision = \frac{Tp}{Tp + Fp} \quad (3)$$

- **Recall** evaluates the model's ability to minimize false negatives by identifying all true positive cases within the dataset, as defined in Equation (4).

$$Recall = \frac{Tp}{Tp + Fn} \quad (4)$$

- **F1-score** computes the harmonic mean of precision and recall, providing a balanced metric that takes into consideration both false positives and false negatives, as demonstrated in Equation (5).

$$F1 - Score = 2 \times \frac{Recall \times Precision}{Recall + Precision} \quad (5)$$

- **AUC (Area Under the ROC Curve)** is a commonly used metric for binary classification models. It represents the area under the ROC (Receiver Operating Characteristic) curve, which displays the True Positive Rate (TPR) versus the False Positive Rate (FPR) across various threshold values [85].

These metrics collectively offer an excellent evaluation of the model's effectiveness and reliability.

4.5. Parameter Setting

The performance of a deep learning model is largely dependent on the hyperparameters' selection. Selecting the appropriate hyperparameters is vital for developing a model that is both accurate and able to generalize effectively. For example, learning rate plays a crucial role during model training. While too low of a learning rate might result in slow convergence and needless training time, too high of a rate can cause early convergence and poor choice of solutions. Furthermore, whereas adding more epochs generally improves model performance, there is a point beyond which further increases may yield diminishing returns. In this study, using the IBM dataset, key parameters such as batch size, dropout rate, number of epochs, loss function, and learning rate have been evaluated. After extensive experimentation and careful consideration of the trade-offs among various parameters, the final configuration of parameters used in the experiments is presented in Table 4.

Table 4. The parameter values used for experiments.

Model	Parameters	Value
All deep learning models	Epoch	50
All deep learning models	Optimizer	Adam (lr = 0.001)
All deep learning models	Loss function	Binary Crossentropy
All deep learning models	Bath Size	128
All deep learning models	Dropout Rate	0.5
CNN	Kernel	3
CNN	# Filters	100
MLP, LSTM, Bi-LSTM, GRU, Bi-GRU	# Units	100
Transformer	# attention heads	8
Transformer	Head-dim	64
Transformer	Feedforward-dim	2048
Bi-TCN 1 (Proposed)	# Filter	32
Bi-TCN 2 (Proposed)	# Filter	64
Bi-TCN 1 (Proposed)	kernel	3
Bi-TCN 2 (Proposed)	kernel	5
Both Bi-TCN (Proposed)	dilations	[1, 2, 4]
Fully connected (Proposed)	# Units	128
RF, AdaBoost, Gradient Boosting	# Estimators	100
Gradient Boosting, CatBoost	Learning rate	1.0
CatBoost	Iterations	10
CatBoost	Depth	2

4.6. Results

This section displays the findings and evaluations of our proposed model on IBM and Kaggle datasets.

4.6.1. Results on IBM Dataset

In order to evaluate our proposed model, we also tested several widely used classical machine learning and deep learning models for comparison. The overall performance of these models on the IBM dataset is summarized in Table 5. The table includes a variety of evaluation measures, including accuracy, precision, recall, F1-score, AUC, and training time, assessed using 5-fold cross-validation. These findings provide insightful information on the positive aspects of the proposed model over alternative strategies, assisting in the identification of the best methods for predicting employee attrition.

Table 5. Performance of the proposed model compared to various models on the IBM dataset.

Models	Accuracy %	Precision %	Recall %	F1-Score %	AUC %	Time (M:S)
Machine Learning Models						
Decision Tree (DT)	77.82	34.39	40.59	36.97	62.81	00:01
Random Forest (RF)	85.98	80.37	17.88	28.95	58.49	00:10
KNN	82.85	40.50	16.54	22.95	56.13	00:01
Logistic Regression	87.68	74.66	36.30	48.58	66.97	00:01
Ada Boost	87.61	72.09	41.27	51.73	68.92	00:11
Gradient Boosting	87.41	66.99	45.34	53.69	70.47	00:06
XG Boost	86.93	70.55	33.73	45.10	65.47	00:03
Cat Boost	87.21	69.50	38.24	49.07	67.45	00:12
Deep Learning Models						
MLP	86.93	69.69	33.04	44.53	65.17	00:15
CNN	87.21	76.81	29.44	42.22	63.90	00:20
LSTM	86.12	58.05	43.15	48.10	68.78	00:24
Bi-LSTM	86.59	58.76	49.44	53.56	71.56	00:33
GRU	86.53	61.28	44.42	51.27	69.48	00:21
BI-GRU	86.87	60.67	52.39	56.15	72.89	00:29
Transformer	88.77	57.78	47.97	51.86	72.23	02:21
Proposed Model (Ensemble Bi-TCN)	89.65	69.81	56.30	61.61	77.58	00:59

According to the results, our proposed model outperforms other baseline deep learning and machine learning models, achieving an accuracy of 89.65% and an F1-score of 61.61%. A closer analysis reveals that the second-best accuracy, following our proposed model, was achieved by Transformer, with an accuracy of 88.77%. In terms of the F1-score, the next best performances were obtained by Bi-GRU and Bi-LSTM, with F1-scores of 56.15% and 53.56%, respectively. In terms of training time, the proposed model takes longer than other models due to its bi-directional and ensemble structure. However, since the data for this specific task are typically small, this drawback can be ignored.

The loss and validation loss diagrams of the proposed model on the IBM dataset are presented in Figure 7a, while Figure 7b illustrates the accuracy and validation accuracy diagrams for the model based on the IBM dataset. The loss diagram illustrates the training loss values, showing how the model's predictions improve and align with the actual labels throughout the training process. The validation loss diagram, on the other hand, depicts the fluctuation in loss values on the validation set during evaluation.

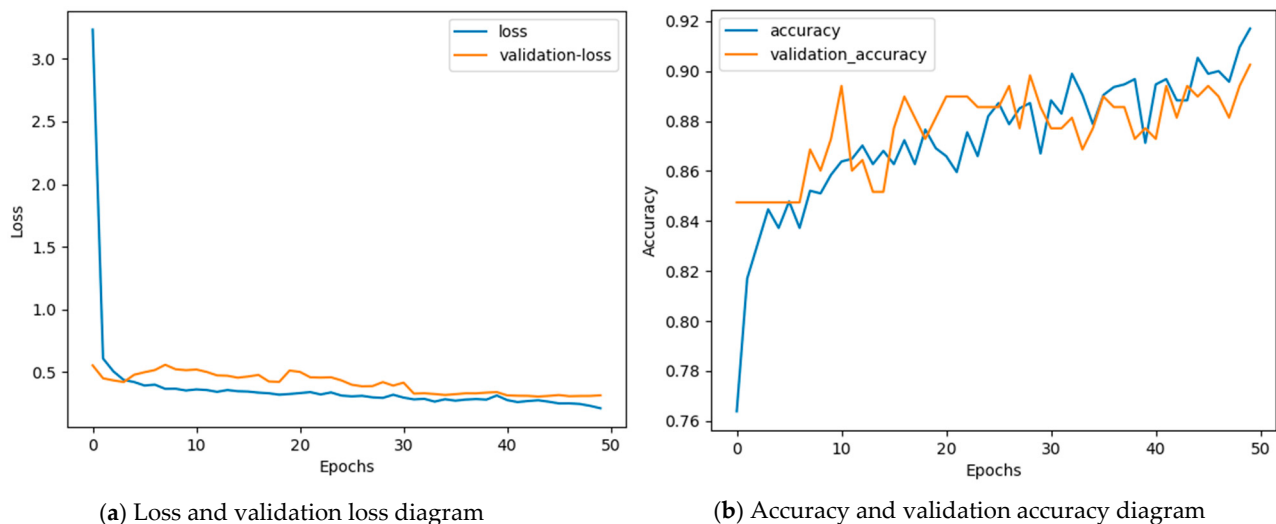


Figure 7. The loss and accuracy diagrams of the proposed model on IBM dataset.

The accuracy diagram visually represents the model's performance during training, highlighting its ability to predict employee attrition. Meanwhile, the validation accuracy diagram reflects the model's generalization ability by tracking accuracy trends on the validation set across multiple epochs. These diagrams provide a thorough understanding of the model's overall performance and optimization.

According to the diagrams, the training loss starts high and drops rapidly within the first few epochs. It then declines gradually with a low slope, reaching its lowest point around the 50th epoch. The validation loss follows a similar trend but exhibits slightly more fluctuations. After the 30th epoch, it becomes more stable. Both losses eventually stabilize at a low value, indicating that the model is learning effectively. Additionally, the validation loss does not significantly diverge from the training loss, suggesting the absence of severe overfitting. The stabilization of loss after a few epochs implies that the model reaches an optimal point without unnecessary fluctuations. Furthermore, both training and validation accuracy generally follow an increasing trend, with some fluctuations. These fluctuations may indicate an imbalance in the data, but the overall upward trend highlights the model's efficiency.

Figure 8 also presents the ROC-AUC (Receiver Operating Characteristic-Area Under Curve) diagram for the proposed model on the IBM dataset. This graph offers insightful information about the model's classification performance, which aids in assessing its efficacy. We can ensure a more precise and successful binary classification method by making well-informed selections regarding threshold modifications and model selection by examining the ROC-AUC curves.

According to the ROC-AUC diagram, the curve rises steeply, reaching a high True Positive Rate (TPR) early, indicating strong performance. A low False Positive Rate (FPR) at the beginning suggests that the model correctly classifies many positive cases before making errors. Additionally, the Area Under the Curve (AUC) is 0.78, demonstrating the model's good discriminative ability.

In additional experiments, to emphasize the critical importance of data augmentation and class balancing, we utilize a GAN-based data augmentation and three oversampling techniques: Random Oversampling, SMOTE, and ADASYN. These techniques are applied to improve the performance of the proposed model. Table 6 compares the results of using these different techniques in conjunction with the proposed model.

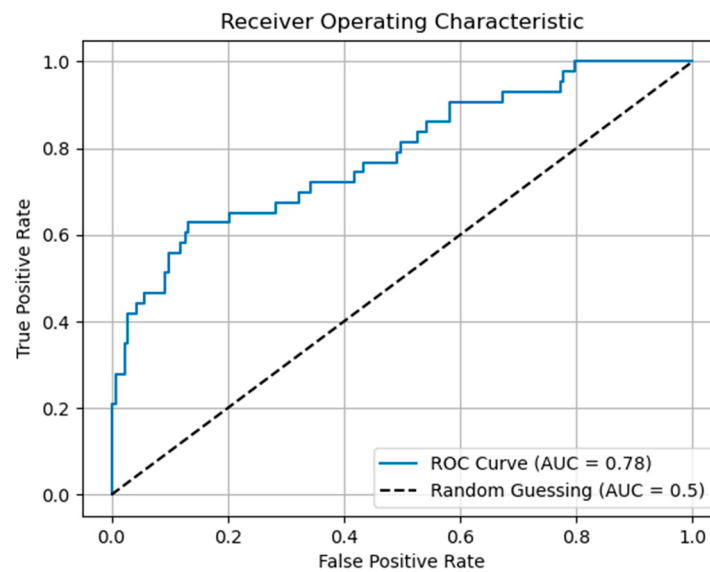


Figure 8. The ROC-AUC diagrams of the proposed model on IBM dataset.

Table 6. Performance of the proposed model with different data augmentation techniques.

Models	Accuracy %	Precision %	Recall %	F1-Score %	AUC %	Time (M:S)
Proposed Model+ GAN	92.17	81.93	66.84	73.60	84.95	03:04
Proposed Model + Random Oversampling	91.83	79.50	74.15	75.99	85.76	01:39
Proposed Model + SMOTE	91.93	90.91	94.11	92.36	91.92	01:40
Proposed Model + ADASYN	92.11	92.33	92.19	92.24	92.11	01:51

For GAN-based data augmentation, we first generated 3000 synthetic data instances based on the original IBM dataset with 1500 belonging to the ‘Yes’ attrition group and 1500 to the ‘No’ attrition group. To ensure the reliability of our results and prevent any bias from synthetic data correlations, we use only the original data for testing. Given that we employ 5-fold cross-validation, in each iteration, 20% of the original data is set aside as test data, while the remaining 80% is combined with the generated synthetic data to form the training set. Consequently, in each iteration, the training set consists of 4176 instances, which is nearly three times larger than the original dataset and effectively addresses the class imbalance by achieving a nearly balanced distribution, while the test set comprises 294 instances. For all three other oversampling methods, the oversampling is applied exclusively to the training data in each iteration.

The results indicate that nearly all techniques produced similar and closely matched outcomes. Notably, all these methods enhanced the efficiency of the proposed model by approximately 2.5%. This improvement can be attributed to the increased volume and improved balance of the training data. However, the model achieved the highest accuracy when using GAN data augmentation, while the best F1-score was obtained with ADASYN.

When comparing these techniques, we can state that GANs produce data that closely mimic genuine samples by learning the entire data distribution and creating completely new synthetic samples for every class. On the other hand, oversampling techniques expand the minority class’s sample size without revealing the distribution as a whole. The Random Oversampling method randomly chooses and duplicates existing instances from the minority class until the class distribution becomes more balanced. SMOTE creates synthetic samples by extrapolating from minority class samples that already exist. Between a randomly chosen data point and one of its K-Nearest Neighbors, a new synthetic sample

is constructed along the line segment. ADASYN takes difficulty and data density into account while creating synthetic data. In areas with limited minority class examples, more synthetic samples are produced, increasing the adaptability of the class distribution.

Figure 9a presents the accuracy and validation accuracy curves for the proposed model with GAN-based data augmentation, while Figure 9b illustrates the loss and validation loss curves for the same model. Additionally, Figure 10 presents the ROC-AUC (Receiver Operating Characteristic–Area Under the Curve) diagram for the proposed model with GAN-based data augmentation on the IBM dataset.

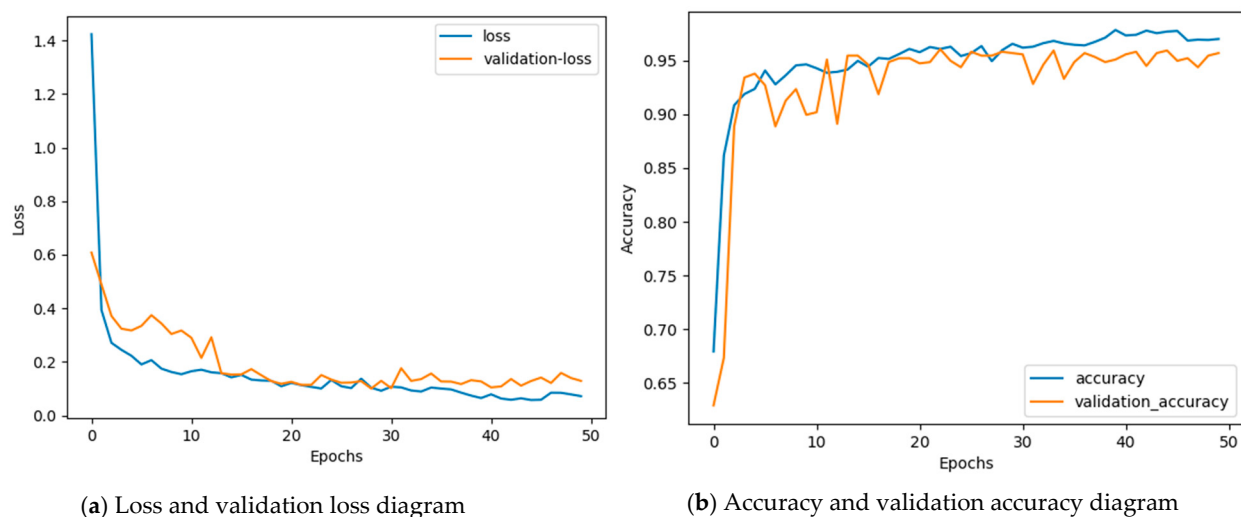


Figure 9. Loss and accuracy diagrams of the proposed model with GAN-based data augmentation.

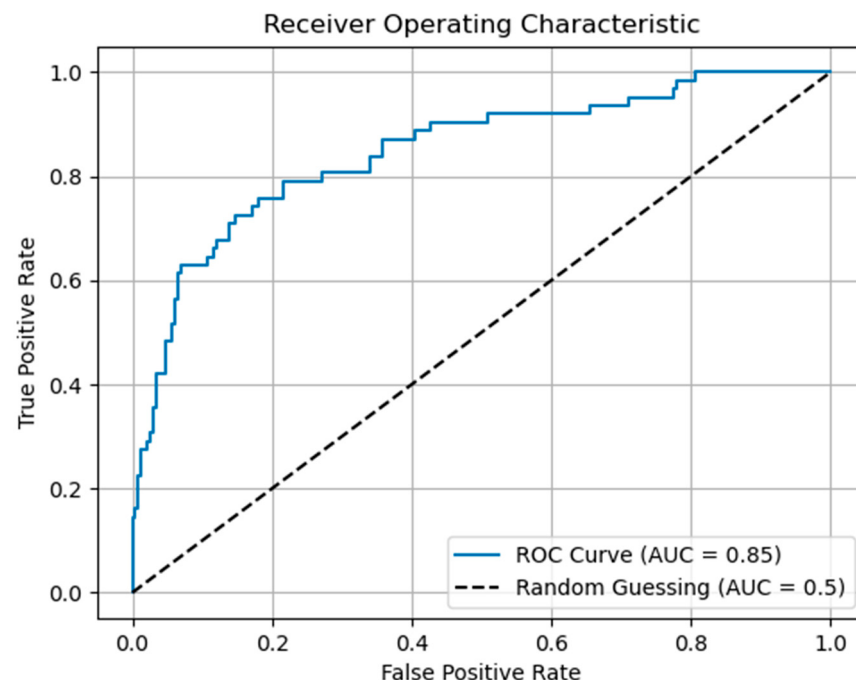


Figure 10. The ROC-AUC diagrams of the proposed model with GAN-based data augmentation.

These diagrams show a clear trend of decreasing loss and increasing accuracy compared to the diagrams for the proposed model without data augmentation. Additionally, they exhibit fewer fluctuations than the diagrams for the model without data augmentation, likely due to the more balanced data.

To identify the key features that have the greatest impact on employee turnover, the explainable AI (XAI) technique, SHAP (Shapley Additive Explanations), was implemented in the proposed model. This approach helps elucidate the features contributing to the model's predictions, providing deeper insights into employee turnover. Figure 11 illustrates the key features in the IBM dataset that have the greatest impact on the prediction of employee turnover.

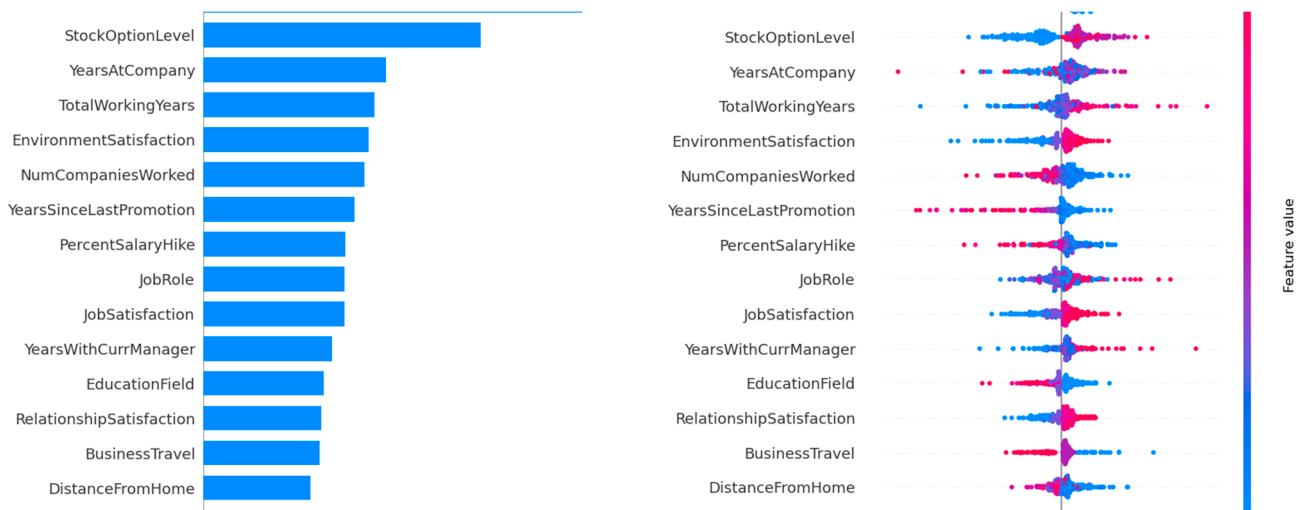


Figure 11. The most influential features in the IBM dataset affecting employee turnover prediction.

SHAP (Shapley Additive Explanations) [39] is an XAI technique that uses a mathematical approach to determine a score for each feature in the model. This score indicates the feature's weight in the model output. Its foundation is game theory, which determines how much each player (feature) contributes to the payout (prediction). In order to determine the scores, it takes into account every possible combination of features to account for both the instances in which the model uses all features and a subset of features. SHAP values provide information on which features are essential for predicting [86].

4.6.2. Results on Kaggle Dataset

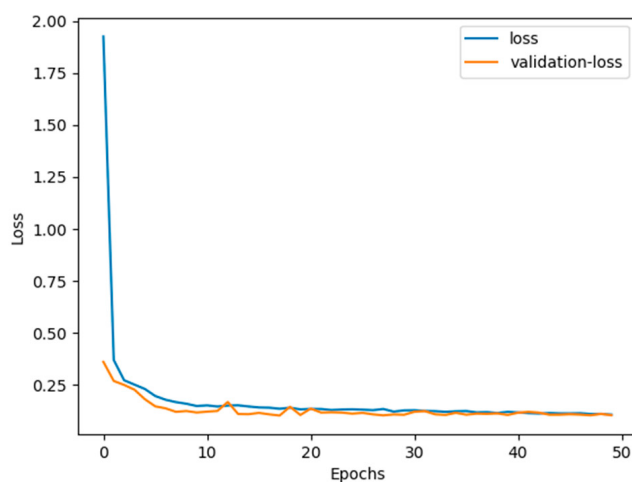
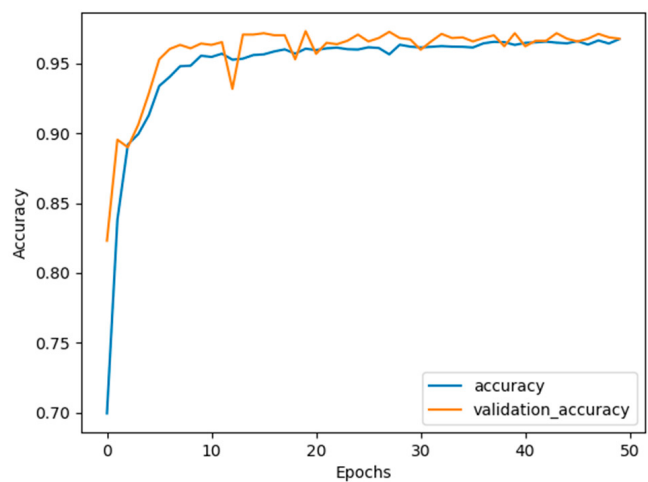
We also evaluate our proposed model on another employee attrition dataset, comparing its performance with baseline classical machine learning and deep learning models. The overall performance of these models on the Kaggle dataset is summarized in Table 7. This table presents various evaluation metrics, including accuracy, precision, recall, F1-score, AUC, and training time, all assessed using 5-fold cross-validation.

The results indicate that the proposed model, RF model, and XGBoost model achieve the best performance on this dataset. The RF model attains the highest accuracy at 97.85%, followed closely by the proposed model with 97.83%, showing only a marginal difference. However, when considering the F1-score metric, the proposed model outperforms the RF model, achieving an F1-score of 95.56% compared to 95.49%. XGBoost ranks next, with an accuracy of 97.62% and an F1-score of 95.10%. Regarding the AUC metric, the proposed model achieves the highest value at 96.94%, followed by the XGBoost and RF models with AUC scores of 96.19% and 96.16%, respectively.

Figure 12a presents the accuracy and validation accuracy diagrams for the proposed model using the Kaggle dataset, while Figure 12b illustrates the loss and validation loss diagrams for the proposed model on the same dataset.

Table 7. Performance of the proposed model compared to various models on the Kaggle dataset.

Models	Accuracy %	Precision %	Recall %	F1-Score %	AUC %	Time (M:S)
Machine Learning Models						
Decision Tree (DT)	95.98	90.82	93.11	91.94	95.01	00:01
Random Forest (RF)	97.85	98.29	92.85	95.49	96.16	00:21
KNN	92.97	82.46	90.77	86.41	92.23	00:02
Logistic Regression	77.64	58.46	31.23	40.66	62.01	00:02
Ada Boost	93.47	86.97	86.35	86.66	91.06	00:21
Gradient Boosting	93.93	88.03	87.28	87.65	91.70	00:17
XG Boost	97.62	96.90	93.38	95.10	96.19	00:04
Cat Boost	94.32	89.33	87.48	88.36	92.02	00:26
Deep Learning Models						
MLP	77.71	58.99	31.24	40.09	62.04	01:34
CNN	77.87	57.83	39.29	45.95	64.84	01:31
LSTM	96.17	94.47	89.81	92.04	94.05	02:02
Bi-LSTM	96.35	93.14	91.99	92.54	94.89	02:29
GRU	96.05	93.80	89.94	91.82	93.99	01:47
BI-GRU	96.55	93.86	92.02	92.93	95.03	02:12
Transformer	96.92	93.31	94.28	93.79	96.03	16:07
Proposed Model (Ensemble Bi-TCN)	97.83	95.95	96.37	95.56	96.94	06:32

**(a)** Loss and validation loss diagram**(b)** Accuracy and validation accuracy diagram**Figure 12.** The loss and accuracy diagrams of the proposed model on Kaggle dataset.

The loss diagrams illustrate that both training and validation loss decrease rapidly, reaching very low values and stabilizing after approximately 20 epochs. This trend suggests that the model is learning efficiently and has converged well. Additionally, the validation loss closely follows the training loss, indicating minimal overfitting.

According to the accuracy diagrams, the training and validation accuracy curves start low and rise rapidly. After 20 epochs, both curves stabilize around 97%, with slight fluctuations in validation accuracy. The close alignment between validation and training accuracy suggests that the model generalizes well.

Figure 13 also presents the ROC-AUC (Receiver Operating Characteristic-Area Under Curve) diagram for the proposed model on the Kaggle dataset.

The steep rise in the ROC curve indicates that the model achieves a high True Positive Rate while maintaining a low False Positive Rate. The dashed line represents random guessing (AUC = 0.5), and since the model's curve is well above this line, it demonstrates

that the proposed model significantly outperforms random guessing. An AUC of 0.97 indicates that the model has a 97% probability of correctly distinguishing between a randomly chosen positive sample and a randomly chosen negative sample, demonstrating excellent performance.

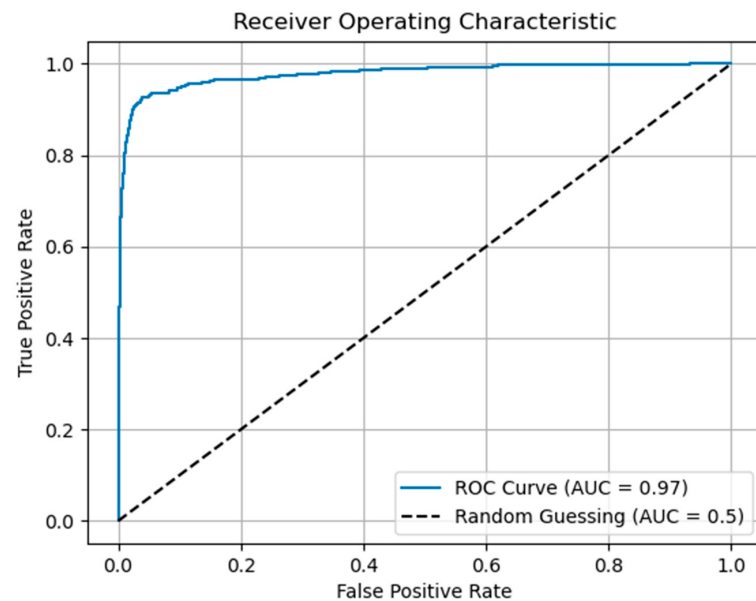


Figure 13. The ROC-AUC diagrams of the proposed model on Kaggle dataset.

We also applied the SHAP (Shapley Additive Explanations) technique in the proposed model to identify the most influential features affecting employee turnover. Figure 14 illustrates the impact of each feature on employee turnover prediction in the Kaggle dataset.

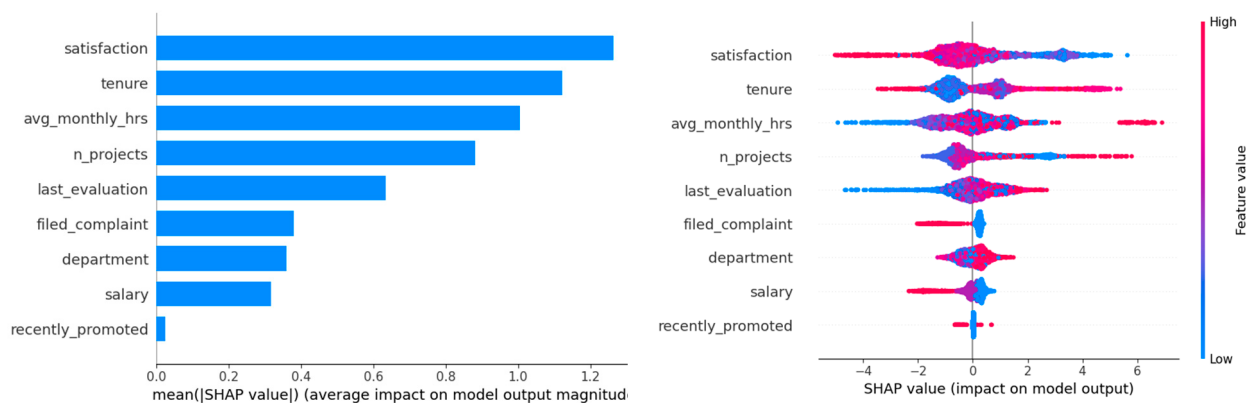


Figure 14. The impact of each feature on employee turnover prediction in the Kaggle dataset.

4.7. Performance Comparison

We compared our method with previous works in the literature that utilized the IBM dataset, as summarized in Table 8. When compared with the most advanced techniques, our approach performed better. Notably, most prior studies focused on classical machine learning models. For instance, the authors of [12] used a Decision Tree model and achieved an accuracy of 82.44%. Logistic Regression (LR) was employed by several researchers, with Mohbey [87] achieving 87% accuracy, the authors of [9] achieving 87.5%, and Qutub et al. [15] improving the accuracy to 88.43% using LR. Another effective model for predicting employee attrition was Support Vector Machine (SVM), which achieved 88.44% accuracy in [88], while the authors of [89] also reported 84% using another SVM variant. Al Akasheh et al. [36] used L-SVM, achieving 87% accuracy, and further improved it to 92.5%

by combining features. Random Forest (RF) was another commonly used machine learning model, achieving accuracies of 80% [35], 85.11% [90], and 87.298% [17]. Boosting algorithms also showed strong performance, with XGBoost achieving 86% [89] and CatBoost achieving 89.45% [26]. Deep learning models have also been explored, with Al-Darraj et al. [30] using a Deep Neural Network (DNN) to achieve an accuracy of 89.11%. Lim et al. [37] introduced a hybrid Genetic Algorithm–Autoencoder–KNN model, achieving an accuracy of 90.95%.

Table 8. Comparing the proposed model with earlier research on the IBM dataset.

Articles	Year	Model	Accuracy
Usha [12]	2019	Decision Tree (DT)48	82.44%
Mohbey [87]	2020	Logistic Regression (LR)	87%
Sethy et al. [88]	2020	SVM	88.44%
Fallucchi et al. [9]	2020	Linear SVC	87.9%
		Gaussian NB	82.5%
		Logistic Regression (LR)	87.5%
		Logistic Regression (LR)	88.43%
Qutub et al. [15]	2021	DT + LR	86.39%
Chakraborty et al. [17]	2021	Random Forest (RF)	87.298%
Najafi-Zangeneh et al. [23]	2021	LR with feature selection	81.0%
		LR without feature selection	78.0%
Al-Darraj et al. [30]	2021	Deep Neural Network (DNN)	89.11%
Mohamed Ahmed [32]	2021	Neural Network	84%
Pratt et al. [90]	2021	Random Forest (RF)	85.11%
Atef et al. [35]	2022	KNN	84%
		RF	80%
Atique et al. [26]	2022	CatBoost	89.45%
Raza et al. [10]	2022	Extra Trees Classifier (10-fold)	93%
Prihanto et al. [89]	2023	XG Boost	86%
		SVM	84%
Al Akasheh et al. [36]	2024	L-SVM	87%
		L-SVM with features combination	92.5%
		Graph Neural Network (GNN)	70.4%
		Graph Attention Network (GAT)	74.3%
Lim et al. [37]	2024	GA-DeepAutoencoder-KNN	90.95%
Proposed Model	2025	Ensemble Bi-TCN	89.65%
Proposed Model with Data Augmentation	2025	Ensemble Bi-TCN+ (GAN)	92.17%

5. Discussion and Conclusions

The primary aim of this study was to support HR managers in mitigating employee attrition by using predictive analytics to identify potential departures as early as possible. This predictive capability enables organizations to save valuable time and resources by reducing recruitment and training efforts. Additionally, it helps businesses meet deadlines and maintain stable staffing levels by preventing turnover. To accomplish this, we introduced a novel deep learning framework for employee attrition prediction based on the Bidirectional Temporal Convolutional Network (Bi-TCN). We utilized two publicly available datasets, IBM and Kaggle, for our experiments. The IBM dataset contained information from 1470 employees with 35 features, while the Kaggle dataset included 14,249 employees and 10 features. We conducted extensive experiments using both datasets, evaluating our model's performance against both cutting-edge methods and conventional machine learning and deep learning models in terms of important metrics including accuracy, precision, recall, F1-score, and AUC. The proposed model achieved an accuracy of 89.65% with five-fold cross-validation on the IBM dataset, and 97.83% on the Kaggle dataset. We also implemented a fully connected GAN-based data augmentation technique along with three oversampling methods including Random Oversampling, SMOTE, and ADASYN, to

enhance and balance the IBM dataset. The results demonstrate that our proposed model, when combined with the GAN-based approach, achieves an accuracy of 92.17%. These results surpass most baseline machine learning and deep learning models, as well as state-of-the-art approaches, highlighting the model's potential to improve current methods and its broad applicability across various industries.

Beyond merely putting algorithms into practice, our study analyzes the results in the particular context of HR decision-making, highlighting the significance of determining important aspects linked to employee attrition. We applied SHAP, an explainable AI method, to determine the most influential features contributing to attrition. By highlighting the crucial factors influencing attrition, the proposed model helps organizations create focused retention and recruitment plans in advance. The cost of recruiting and onboarding new staff can be significantly reduced with this calculated strategy. Actually, by utilizing cutting-edge deep learning models, data augmentation strategies, and explainability methodologies, this research seeks to provide HR managers with deeper insights so they may make better-informed and efficient decisions about employee retention. By knowing what causes turnover, HR practices may be improved, which will increase employee retention and create a more effective workplace. The findings can also be used to evaluate how well a worker's abilities, values, and traits match the demands of the position. In addition to improving turnover prediction, this method offers insightful information for maximizing employee–job fit, which eventually fortifies retention tactics.

Despite the promising results, we acknowledge the limitations of our approach. The current model evaluation is based only on two datasets, which are useful but may not fully reflect the complexities of employee attrition in different organizations. Future research should focus on validating the model's generalizability by incorporating a wider range of datasets from various industries and organizational structures. This will ensure the robustness and adaptability of the proposed approach in real-world business environments. Moreover, while our model effectively leverages structured data for employee turnover prediction, it does not explicitly account for psychological and subjective factors that often influence attrition decisions. Employees' sentiments, motivations, and perceptions play a crucial role in determining their likelihood of retention or departure. Future studies should explore the integration of advanced psychological assessments, sentiment analysis from employee feedback, and surveys measuring job satisfaction, well-being, and workplace engagement. Employing techniques such as natural language processing (NLP) on employee reviews and interviews, combined with psychometric evaluations, could provide deeper insights into the underlying causes of attrition. By incorporating these nuanced, human-centric factors into predictive frameworks, organizations can develop more comprehensive and accurate models for workforce analytics.

Author Contributions: Conceptualization, F.M.S. and S.Y.; methodology, F.M.S.; software, F.M.S.; validation S.Y. and M.A.B.A.; investigation, F.M.S. and M.A.B.A.; writing—original draft preparation, F.M.S.; writing—review and editing, S.Y. and M.A.B.A.; supervision, S.Y.; project administration, S.Y.; funding acquisition, S.Y. All authors have read and agreed to the published version of the manuscript.

Funding: This research received no external funding.

Institutional Review Board Statement: Not applicable.

Informed Consent Statement: Not applicable.

Data Availability Statement: Publicly available datasets were analyzed in this study. These datasets can be found here: <https://www.kaggle.com/datasets/pavansubhasht/ibm-hr-analytics-attrition-dataset> (accessed on 1 February 2025) and <https://www.kaggle.com/code/devisangeetha/predicting-employee-churn> (accessed on 1 February 2025).

Conflicts of Interest: The authors declare no conflicts of interest.

Abbreviations

The following abbreviations are used in this manuscript:

Bi-TCN	Bidirectional Temporal Convolutional Network
DT	Decision Tree
KNN	K-Nearest Neighbors
RF	Random Forest
LR	Logistic Regression
AdaBoost	Adaptive Boosting
GB	Gradient Boosting
XGBoost	Extreme Gradient Boosting
CatBoost	Categorical Boosting
CNN	Convolutional Neural Network
MLP	Multi-Layer Perceptron
LSTM	Long Short-Term Memory
GRU	Gated Recurrent Unit
Bi-LSTM	Bidirectional Long Short-Term Memory
Bi-GRU	Bidirectional Gated Recurrent Unit

References

1. Vardarlier, P.; Zafer, C. Use of artificial intelligence as business strategy in recruitment process and social perspective. In *Digital Business Strategies in Blockchain Ecosystems: Transformational Design and Future of Global Business*; Springer: Cham, Switzerland, 2020; pp. 355–373. [\[CrossRef\]](#)
2. Al Akasheh, M.; Malik, E.F.; Hujran, O.; Zaki, N. A decade of research on machine learning techniques for predicting employee turnover: A systematic literature review. *Expert Syst. Appl.* **2024**, *238*, 121794. [\[CrossRef\]](#)
3. Alduayj, S.S.; Rajpoot, K. Predicting employee attrition using machine learning. In Proceedings of the 2018 International Conference on Innovations in Information Technology (iit), Al Ain, United Arab Emirates, 18–19 November 2018; pp. 93–98.
4. Alsheref, F.K.; Fattoh, I.E.; Ead, W.M. Automated prediction of employee attrition using ensemble model based on machine learning algorithms. *Comput. Intell. Neurosci.* **2022**, *2022*, 7728668. [\[CrossRef\]](#) [\[PubMed\]](#)
5. Setiawan, I.A.; Suprihanto, S.; Nugraha, A.; Hutahaeen, J. HR analytics: Employee attrition analysis using logistic regression. *IOP Conf. Ser. Mater. Sci. Eng.* **2020**, *830*, 032001. [\[CrossRef\]](#)
6. Jain, P.K.; Jain, M.; Pamula, R. Explaining and predicting employees' attrition: A machine learning approach. *SN Appl. Sci.* **2020**, *2*, 757. [\[CrossRef\]](#)
7. Rane, N.L.; Paramesha, M.; Choudhary, S.P.; Rane, J. Artificial intelligence, machine learning, and deep learning for advanced business strategies: A review. *Partn. Univers. Int. Innov. J.* **2024**, *2*, 147–171. [\[CrossRef\]](#)
8. Loureiro, S.M.C.; Guerreiro, J.; Tussyadiah, I. Artificial intelligence in business: State of the art and future research agenda. *J. Bus. Res.* **2021**, *129*, 911–926. [\[CrossRef\]](#)
9. Fallucchi, F.; Coladangelo, M.; Giuliano, R.; William De Luca, E. Predicting employee attrition using machine learning techniques. *Computers* **2020**, *9*, 86–103. [\[CrossRef\]](#)
10. Raza, A.; Munir, K.; Almutairi, M.; Younas, F.; Fareed, M.M.S. Predicting employee attrition using machine learning approaches. *Appl. Sci.* **2022**, *12*, 6424. [\[CrossRef\]](#)
11. Ajit, P. Prediction of employee turnover in organizations using machine learning algorithms. *(IJARAI) Int. J. Adv. Res. Artif. Intell.* **2016**, *5*, 22–25.
12. PM, U.; Balaji, N. Analysing Employee attrition using machine learning. *Karpagam J. Comput. Sci.* **2019**, *13*, 277–282.
13. Guerranti, F.; Dimitri, G.M. A comparison of machine learning approaches for predicting employee attrition. *Appl. Sci.* **2022**, *13*, 267. [\[CrossRef\]](#)
14. Zhao, Y.; Hryniewicki, M.K.; Cheng, F.; Fu, B.; Zhu, X. Employee turnover prediction with machine learning: A reliable approach. In *Intelligent Systems and Applications: Proceedings of the 2018 Intelligent Systems Conference (IntelliSys)*; Springer: Cham, Switzerland, 2019; Volume 2, pp. 737–758.
15. Qutub, A.; Al-Mehmadi, A.; Al-Hssan, M.; Aljohani, R.; Alghamdi, H.S. Prediction of employee attrition using machine learning and ensemble methods. *Int. J. Mach. Learn. Comput* **2021**, *11*, 110–114. [\[CrossRef\]](#)

16. El-Rayes, N.; Fang, M.; Smith, M.; Taylor, S.M. Predicting employee attrition using tree-based models. *Int. J. Organ. Anal.* **2020**, *28*, 1273–1291. [\[CrossRef\]](#)
17. Chakraborty, R.; Mridha, K.; Shaw, R.N.; Ghosh, A. Study and prediction analysis of the employee turnover using machine learning approaches. In Proceedings of the 2021 IEEE 4th International Conference on Computing, Power and Communication Technologies (GUCON), Kuala Lumpur, Malaysia, 24–26 September 2021; pp. 1–6.
18. Alzate Vanegas, J.M.; Wine, W.; Drasgow, F. Predictions of attrition among US Marine Corps: Comparison of four predictive methods. *Mil. Psychol.* **2022**, *34*, 147–166. [\[CrossRef\]](#)
19. Gao, X.; Wen, J.; Zhang, C. An improved random forest algorithm for predicting employee turnover. *Math. Probl. Eng.* **2019**, *2019*, 4140707. [\[CrossRef\]](#)
20. Jain, N.; Jana, P.K. Xrrf: An explainable reasonably randomised forest algorithm for classification and regression problems. *Inf. Sci.* **2022**, *613*, 139–160. [\[CrossRef\]](#)
21. Wild Ali, A.B. Prediction of employee turn over using random forest classifier with intensive optimized PCA algorithm. *Wirel. Pers. Commun.* **2021**, *119*, 3365–3382. [\[CrossRef\]](#)
22. Ponnuru, S.a.; Merugumala, G.; Padigala, S.; Vanga, R.; Kantapalli, B. Employee attrition prediction using logistic regression. *Int. J. Res. Appl. Sci. Eng. Technol.* **2020**, *8*, 2871–2875. [\[CrossRef\]](#)
23. Najafi-Zangeneh, S.; Shams-Gharneh, N.; Arjomandi-Nezhad, A.; Hashemkhani Zolfani, S. An improved machine learning-based employees attrition prediction framework with emphasis on feature selection. *Mathematics* **2021**, *9*, 1226. [\[CrossRef\]](#)
24. Jain, R.; Nayyar, A. Predicting employee attrition using xgboost machine learning approach. In Proceedings of the 2018 International Conference on System Modeling & Advancement in Research Trends (Smart), Moradabad, India, 23–24 November 2018; pp. 113–120.
25. Jhaver, M.; Gupta, Y.; Mishra, A.K. Employee turnover prediction system. In Proceedings of the 2019 4th International Conference on Information Systems and Computer Networks (ISCON), Mathura, India, 21–22 November 2019; pp. 391–394.
26. Atique, M.M.A.B.; Hoque, M.N.; Uddin, M.J. Employee Attrition Analysis Using CatBoost. In Proceedings of the International Conference on Machine Intelligence and Emerging Technologies, Noakhali, Bangladesh, 23–25 September 2022; pp. 644–658.
27. Jain, N.; Tomar, A.; Jana, P.K. A novel scheme for employee churn problem using multi-attribute decision making approach and machine learning. *J. Intell. Inf. Syst.* **2021**, *56*, 279–302. [\[CrossRef\]](#)
28. Yahia, N.B.; Hlel, J.; Colomo-Palacios, R. From big data to deep data to support people analytics for employee attrition prediction. *IEEE Access* **2021**, *9*, 60447–60458. [\[CrossRef\]](#)
29. Pekel Ozmen, E.; Ozcan, T. A novel deep learning model based on convolutional neural networks for employee churn prediction. *J. Forecast.* **2022**, *41*, 539–550. [\[CrossRef\]](#)
30. Al-Darraj, S.; Honi, D.G.; Fallucchi, F.; Abdulsada, A.I.; Giuliano, R.; Abdulmalik, H.A. Employee attrition prediction using deep neural networks. *Computers* **2021**, *10*, 141. [\[CrossRef\]](#)
31. Brown, A.; Davis, N.; Miller, O.; Wilson, E.; Smith, L.; Lopez, S. Deep Learning Techniques for Enhancing Employee Turnover Prediction Accuracy. **2024**. [\[CrossRef\]](#)
32. Mohamed Ahmed, T. A novel classification model for employees turnover using neural network to enhance job satisfaction in organizations. *J. Inf. Organ. Sci.* **2021**, *45*, 361–374. [\[CrossRef\]](#)
33. Zhu, Q.; Shang, J.; Cai, X.; Jiang, L.; Liu, F.; Qiang, B. CoxRF: Employee turnover prediction based on survival analysis. In Proceedings of the 2019 IEEE SmartWorld, Ubiquitous Intelligence & Computing, Advanced & Trusted Computing, Scalable Computing & Communications, Cloud & Big Data Computing, Internet of People and Smart City Innovation (SmartWorld/SCALCOM/UIC/ATC/CBDCom/IOP/SCI), Leicester, UK, 19–23 August 2019; pp. 1123–1130.
34. Jin, Z.; Shang, J.; Zhu, Q.; Ling, C.; Xie, W.; Qiang, B. RFRSF: Employee turnover prediction based on random forests and survival analysis. In Proceedings of the Web Information Systems Engineering–WISE 2020: 21st International Conference, Amsterdam, The Netherlands, 20–24 October 2020; Proceedings, Part II 21. pp. 503–515.
35. Atef, M.; S Elzanfaly, D.; Ouf, S. Early prediction of employee turnover using machine learning algorithms. *Int. J. Electr. Comput. Eng. Syst.* **2022**, *13*, 135–144. [\[CrossRef\]](#)
36. Al Akasheh, M.; Hujran, O.; Malik, E.F.; Zaki, N. Enhancing the Prediction of Employee Turnover with Knowledge Graphs and Explainable AI. *IEEE Access* **2024**, *12*, 77041–77053. [\[CrossRef\]](#)
37. Lim, C.S.; Malik, E.F.; Khaw, K.W.; Alnoor, A.; Chew, X.; Chong, Z.L.; Al Akasheh, M. Hybrid GA–DeepAutoencoder–KNN Model for Employee Turnover Prediction. *Stat. Optim. Inf. Comput.* **2024**, *12*, 75–90. [\[CrossRef\]](#)
38. Marín Díaz, G.; Galán Hernández, J.J.; Galdón Salvador, J.L. Analyzing employee attrition using explainable AI for strategic HR decision-making. *Mathematics* **2023**, *11*, 4677. [\[CrossRef\]](#)
39. Lundberg, S.M.; Lee, S.-I. A unified approach to interpreting model predictions. *Adv. Neural Inf. Process. Syst.* **2017**, *30*, 4765–4774.
40. Ribeiro, M.T.; Singh, S.; Guestrin, C. “Why should I trust you?” Explaining the predictions of any classifier. In Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining, San Francisco, CA, USA, 13–17 August 2016; pp. 1135–1144.

41. Makanga, C.; Mukwaba, D.; Agaba, C.L.; Murindanyi, S.; Joseph, T.; Hellen, N.; Marvin, G. Explainable Machine Learning and Graph Neural Network Approaches for Predicting Employee Attrition. In Proceedings of the 2024 Sixteenth International Conference on Contemporary Computing, Noida, India, 8–10 August 2024; pp. 243–255.
42. Shafie, M.R.; Khosravi, H.; Farhadpour, S.; Das, S.; Ahmed, I. A cluster-based human resources analytics for predicting employee turnover using optimized Artificial Neural Networks and data augmentation. *Decis. Anal. J.* **2024**, *11*, 100461. [\[CrossRef\]](#)
43. Varkiani, S.M.; Pattarin, F.; Fabbri, T.; Fantoni, G. Predicting employee attrition and explaining its determinants. *Expert Syst. Appl.* **2025**, *272*, 126575. [\[CrossRef\]](#)
44. Lunardon, N.; Menardi, G.; Torelli, N. ROSE: A package for binary imbalanced learning. *R J.* **2014**, *6*, 79–89. [\[CrossRef\]](#)
45. Shiri, F.M.; Perumal, T.; Mustapha, N.; Mohamed, R.; Ahmadon, M.A.B.; Yamaguchi, S. Measuring Student Satisfaction Based on Analysis of Physical Parameters in Smart Classroom. In Proceedings of the 2024 12th International Conference on Information and Education Technology (ICIET), Yamaguchi, Japan, 18–20 March 2024; pp. 18–23.
46. Zhang, G.; Wang, C.; Xu, B.; Grosse, R. Three mechanisms of weight decay regularization. *arXiv* **2018**, arXiv:1810.12281.
47. Lea, C.; Flynn, M.D.; Vidal, R.; Reiter, A.; Hager, G.D. Temporal convolutional networks for action segmentation and detection. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR), Honolulu, HI, USA, 21–26 July 2017; pp. 156–165.
48. Shiri, F.M.; Perumal, T.; Mustapha, N.; Mohamed, R. A Comprehensive Overview and Comparative Analysis on Deep Learning Models. *J. Artif. Intell.* **2024**, *6*, 301–360. [\[CrossRef\]](#)
49. Zhu, J.; Su, L.; Li, Y. Wind power forecasting based on new hybrid model with TCN residual modification. *Energy AI* **2022**, *10*, 100199. [\[CrossRef\]](#)
50. He, K.; Zhang, X.; Ren, S.; Sun, J. Deep residual learning for image recognition. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR), Las Vegas, NV, USA, 27–30 June 2016; pp. 770–778.
51. Zhang, X.; Dong, F.; Chen, G.; Dai, Z. Advance prediction of coastal groundwater levels with temporal convolutional and long short-term memory networks. *Hydrol. Earth Syst. Sci.* **2023**, *27*, 83–96. [\[CrossRef\]](#)
52. Subhash, P.; IBM HR Analytics Employee Attrition & Performance. Kaggle. 2017. Available online: <https://www.kaggle.com/datasets/pavansubhasht/ibm-hr-analytics-attrition-dataset> (accessed on 6 March 2025).
53. Mzinic. Employee Churn Prediction. Kaggle. 2019. Available online: <https://www.kaggle.com/code/devisangeetha/predicting-employee-churn> (accessed on 6 March 2025).
54. Alzubaidi, L.; Zhang, J.; Humaidi, A.J.; Al-Dujaili, A.; Duan, Y.; Al-Shamma, O.; Santamaria, J.; Fadhel, M.A.; Al-Amidie, M.; Farhan, L. Review of deep learning: Concepts, CNN architectures, challenges, applications, future directions. *J. Big Data* **2021**, *8*, 1–74. [\[CrossRef\]](#)
55. Singh, D.; Merdivan, E.; Kropf, J.; Holzinger, A. Class imbalance in multi-resident activity recognition: An evaluative study on explainability of deep learning approaches. *Univers. Access. Inf. Soc.* **2024**, 1–19. [\[CrossRef\]](#)
56. Iwana, B.K.; Uchida, S. An empirical survey of data augmentation for time series classification with neural networks. *PLoS ONE* **2021**, *16*, e0254841. [\[CrossRef\]](#)
57. Khosla, C.; Saini, B.S. Enhancing performance of deep learning models with different data augmentation techniques: A survey. In Proceedings of the 2020 International Conference on Intelligent Engineering and Management (ICIEM), London, UK, 17–19 June 2020; pp. 79–85.
58. Jimale, A.O.; Noor, M.H.M. Fully connected generative adversarial network for human activity recognition. *IEEE Access* **2022**, *10*, 100257–100266. [\[CrossRef\]](#)
59. Goodfellow, I.; Pouget-Abadie, J.; Mirza, M.; Xu, B.; Warde-Farley, D.; Ozair, S.; Courville, A.; Bengio, Y. Generative adversarial nets. *Adv. Neural. Inf. Process. Syst.* **2014**, *27*, 2672–2680.
60. Kwon, H. Untargeted Evasion Attacks on Deep Neural Networks Using StyleGAN. *Electronics* **2025**, *14*, 574. [\[CrossRef\]](#)
61. Wang, C.; Xu, C.; Yao, X.; Tao, D. Evolutionary generative adversarial networks. *IEEE Trans. Evol. Comput.* **2019**, *23*, 921–934. [\[CrossRef\]](#)
62. Chan, M.H.; Noor, M.H.M. A unified generative model using generative adversarial network for activity recognition. *J. Ambient Intell. Humaniz. Comput.* **2021**, *12*, 8119–8128. [\[CrossRef\]](#)
63. Chawla, N.V.; Bowyer, K.W.; Hall, L.O.; Kegelmeyer, W.P. SMOTE: Synthetic minority over-sampling technique. *J. Artif. Intell. Res.* **2002**, *16*, 321–357. [\[CrossRef\]](#)
64. He, H.; Bai, Y.; Garcia, E.A.; Li, S. ADASYN: Adaptive synthetic sampling approach for imbalanced learning. In Proceedings of the 2008 IEEE International Joint Conference on Neural Networks, Hong Kong, China, 1–8 June 2008; pp. 1322–1328.
65. Mukherjee, M.; Khushi, M. SMOTE-ENC: A novel SMOTE-based method to generate synthetic data for nominal and continuous features. *Appl. Syst. Innov.* **2021**, *4*, 18. [\[CrossRef\]](#)
66. Figueira, A.; Vaz, B. Survey on synthetic data generation, evaluation methods and GANs. *Mathematics* **2022**, *10*, 2733. [\[CrossRef\]](#)
67. Abdelrahman, O.; Keikhosrokiani, P. Assembly line anomaly detection and root cause analysis using machine learning. *IEEE Access* **2020**, *8*, 189661–189672. [\[CrossRef\]](#)

68. Breiman, L. Random forests. *Mach. Learn.* **2001**, *45*, 5–32. [CrossRef]
69. Shiri, F.M.P.; Perumal, T.; Mustapha, N.; Mohamed, R.; Ahmadon, M.A.; Yamaguchi, S. A Survey on Multi-Resident Activity Recognition in Smart Environments. In *Evolution of Information, Communication and Computing System*; UTHM: Parit Raja, Malaysia, 2023; pp. 12–27.
70. Freund, Y.; Schapire, R.E. Experiments with a new boosting algorithm. In Proceedings of the ICML, Bari, Italy, 3–6 July 1996; pp. 148–156.
71. Sarker, I.H. Machine learning: Algorithms, real-world applications and research directions. *SN Comput. Sci.* **2021**, *2*, 160. [CrossRef]
72. Friedman, J.; Hastie, T.; Tibshirani, R. Additive logistic regression: A statistical view of boosting. *Ann. Stat.* **2000**, *28*, 337–407. [CrossRef]
73. Guillen, M.D.; Aparicio, J.; Esteve, M. Gradient tree boosting and the estimation of production frontiers. *Expert Syst. Appl.* **2023**, *214*, 119134. [CrossRef]
74. Prokhorenkova, L.; Gusev, G.; Vorobev, A.; Dorogush, A.V.; Gulin, A. CatBoost: Unbiased boosting with categorical features. In *Advances in Neural Information Processing Systems*; 2018; Volume 31, Available online: https://proceedings.neurips.cc/paper_files/paper/2018/file/14491b756b3a51daac41c24863285549-Paper.pdf (accessed on 6 March 2025).
75. Sarker, I.H. Deep learning: A comprehensive overview on techniques, taxonomy, applications and research directions. *SN Comput. Sci.* **2021**, *2*, 420. [CrossRef] [PubMed]
76. Li, Z.; Liu, F.; Yang, W.; Peng, S.; Zhou, J. A Survey of Convolutional Neural Networks: Analysis, Applications, and Prospects. *IEEE Trans. Neural Netw. Learn. Syst.* **2022**, *33*, 6999–7019. [CrossRef]
77. Babu, B.P.; Narayanan, S.J. One-vs-All Convolutional Neural Networks for Synthetic Aperture Radar Target Recognition. *Cybern. Inf. Technol.* **2022**, *22*, 179–197. [CrossRef]
78. Graves, A. Generating sequences with recurrent neural networks. *arXiv* **2013**, arXiv:1308.0850.
79. Shiri, F.M.; Ahmadi, E.; Rezaee, M.; Perumal, T. Detection of Student Engagement in E-Learning Environments Using EfficientnetV2-L Together with RNN-Based Models. *J. Artif. Intell.* **2024**, *6*, 85–103. [CrossRef]
80. Chung, J.; Gulcehre, C.; Cho, K.; Bengio, Y. Empirical evaluation of gated recurrent neural networks on sequence modeling. *arXiv* **2014**, arXiv:1412.3555.
81. Shiri, F.M.; Perumal, T.; Mustapha, N.; Mohamed, R.; Ahmadon, M.A.B.; Yamaguchi, S. Recognition of Student Engagement and Affective States Using ConvNeXtLarge and Ensemble GRU in E-Learning. In Proceedings of the 2024 12th International Conference on Information and Education Technology (ICIET), Yamaguchi, Japan, 18–20 March 2024; pp. 30–34.
82. Vaswani, A.; Shazeer, N.; Parmar, N.; Uszkoreit, J.; Jones, L.; Gomez, A.N.; Kaiser, Ł.; Polosukhin, I. Attention is all you need. In Proceedings of the 31st International Conference on Neural Information Processing Systems, Long Beach, CA, USA, 4–9 December 2017; pp. 5998–6008.
83. Zhang, X.; Liu, C.-A. Model averaging prediction by K-fold cross-validation. *J. Econom.* **2023**, *235*, 280–301. [CrossRef]
84. Aprihartha, M.A.; Idham, I. Optimization of Classification Algorithms Performance with k-Fold Cross Validation. *Eig. Math. J.* **2024**, *7*, 61–66. [CrossRef]
85. Nellore, S.B. Various performance measures in Binary classification-An Overview of ROC study. *IJISSET—Int. J. Innov. Sci. Eng. Technol.* **2015**, *2*, 596–605.
86. Salih, A.M.; Raisi-Estabragh, Z.; Galazzo, I.B.; Radeva, P.; Petersen, S.E.; Lekadir, K.; Menegaz, G. A perspective on explainable artificial intelligence methods: SHAP and LIME. *Adv. Intell. Syst.* **2025**, *7*, 2400304. [CrossRef]
87. Mohbey, K.K. Employee's attrition prediction using machine learning approaches. In *Machine Learning and Deep Learning in Real-Time Applications*; IGI Global: New York, NY, USA, 2020; pp. 121–128.
88. Sethy, A.; Raut, A.K. Employee attrition rate prediction using machine learning approach. *Turk. J. Physiother. Rehabil.* **2020**, *32*, 14024–14031.
89. Prihanto, B.; Sereati, C.O.; Kartawidjaja, M.A.; Siregar, M. Attrition Analysis using XG Boost and Support Vector Machine Algorithm. *Int. J. Innov. Sci. Res. Technol.* **2023**, *8*, 2096–2112.
90. Pratt, M.; Boudhane, M.; Cakula, S. Employee attrition estimation using random forest algorithm. *Balt. J. Mod. Comput.* **2021**, *9*, 49–66. [CrossRef]

Disclaimer/Publisher's Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.