

# Hybrid Deep Learning Model for Diabetes Prediction and Automated Diet Recommendation Using Ensemble Classifiers and Optimized Algorithms

Anjali Jain<sup>\*1</sup>, Alka Singhal<sup>2</sup>

<sup>1</sup>Jaypee Institute of Information Technology, Noida, India & Krishna Institute of Engineering & Technology (KIET), Ghaziabad, Delhi-NCR, India

<sup>2</sup>Jaypee Institute of Information Technology, Noida, India

**Abstract:** Nowadays, Non-Insulin-Dependent Diabetes Mellitus (NIDDM) is one of the most life-threatening and multifaceted illnesses. Many individuals are afflicted with numerous ailments because of insufficient nutrition for diabetic patients. Numerous machine learning practices and approaches adopted and suggested by researchers for predicting diabetes in the current past. In this paper, various classifiers including ensemble voting classifier, artificial neural network, AdaBoost, Random Forest, Support Vector Machine, Gradient Boost and deep learning models like Tabnet, XgBoost, is used to detect diabetic patients from reported dataset and recommend the required calories for a diabetic patient in the form of diet. Further, the results were equated with the prevailing models in terms of classification and recommendation. The PIMA India diabetes dataset is considered to check the performance of the model, which reveal that the Voting Classifier achieved a remarkable accuracy rate of 98.85%. The optimized hybrid deep learning model is designed to perform the recommendation process that helps to select the meals in breakfast, lunch and dinner for diabetes patients. Various performance metrics such as recall, F1 score, accuracy and precision are analyzed on the proposed method. This paper proposed a novel automated Diet Recommendation to recommend the nutrition diet plan for diabetic patients depending on patient's health and medical records. The diet-recommendation system is examined using a hybrid deep learning model called Triple Attention-based Gated Stacked LSTM (TriAtt-St-LSTM). Conclusion: The proposed work has achieved a highest accuracy of 98.09% for the UCI dataset and 99.45% for the diet recommendation system dataset.

**Keywords:** Diet Recommendation, Ensemble Voting classifier, LSTM, PIMA Dataset, XgBoost.

## 1. INTRODUCTION

In today's lifestyle, people are so busy and frequently skipping their meals. Nearly 90% of them require a diet plan which is healthy and gives sufficient energy [1]. Lot of factors impacts the health of an individual like sleep, nutrition, heredity etc. Most of the scenarios of diabetes require early prediction, intervention, and timely control. However, it can be identified, managed, and avoided through proper diet, exercise, medication, routine screenings, and treatment for its side effects. Several factors, such as health status, nutritional needs, dietary access, existing diseases, and technological advancements, play a substantial role in shaping an individual's dietary choices.

Various researchers given their ideas, models, and online tools for diabetes prediction, risk assessment, and tailored dietary recommendations, seeking to actively manage their condition through these diverse digital avenues [2]. A lot of research was conducted by various authors to make sure that diet which contains appropriate proportion of carbohydrates, fats, proteins, vitamins, minerals, and sugar etc. is the balanced diet [3]. One must follow and stick to a nutritious balanced diet plan with required number of recommended calories. Various unprecedented changes in climate, environment and lifestyle of human life, various diseases which used to occur from hierarchy now occur due to lack of nutrition and biological dysfunction. Almost every bit of the population wants to eat healthy and being healthy, which requires lot of effort and schedule organization [4] [5].

Noticeably, various factors are considered while foreseeing if a person is having diabetes mellitus or not. The factors which contribute to identifying the diabetes mellitus can age, insulin, glucose BMI, skin thickness and diabetes pedigree function [6]. Diabetic diets aren't one-size-fits-all. Calorie necessities and nutritional desires vary, so personalized recommendations are essential for effective management [7]. Nowadays, the demand of diet recommendation has increased to the level which leads to the generation of so many diet plans like Low Carbs, Low Fat, Keto, Paleo and Vegan Diet etc. [8]. Every diet plan has its own benefit like weight loss, weight gain or to remain fit. Diabetic people must follow a strict diet and routine with specific calorie count.

Based on the sensitivity of life threatening disease like diabetic and the utilization of the potential of ML approaches for prediction, and for recommendation this paper considers weight and height of a person i.e. Body Mass Index (BMI) and Basal Metabolic Rate (BMR) for evaluating underweight, healthy, and overweight categories and proposed a Healthy Diet Recommendation.

The work drift of the paper is defined as:

*First:* In proposed work, the hybrid ensemble classifier is introduced which tests on Mendeley PIMA Indian Diabetes Dataset: This dataset carries the 768 instances and 8 features, offering a valuable snapshot of Indian diabetic patients (<https://data.mendeley.com/datasets/7zcc8v6hvp/1>) [9].

*Second:* The proposed ensemble considers the linear and nonlinear features of reported datasets. Feature engineering is performed to achieve superior performance over its counterparts, considered in this study.

*Third:* A novel Healthy Diet Recommendation is designed to recommend a nutrition diet plan for diabetic patients depending on patient's health and medical records and after analyzing various parameters. The diet-recommended system is examined using a hybrid deep learning model called Triple Attention-based Gated Stacked LSTM (TriAtt-St-LSTM).

## 2. RELATED WORK

For the diabetes diagnosis, a variety of research has been done so far. Healthcare data related to diabetes can be collected using various methods and analyzed using Deep Learning and Machine Learning (ML) techniques. As it's proved by various other authors, type-2 diabetes is more hazardous, and can be cured at its early stage. To address the type-2 issue A. Mohebbi et al. [5] worked on Convolutional Neural Network (CNN) aimed at type 2 diabetes detection compared with Linear Regression (LR) and Multi-Layer Perceptron (MLP). The accuracy of this model achieved 77.5% in area under curve.

Li. et al. [6] projected a method which suggests the potential of AI in enhancing early disease diagnosis and reducing the risk of infections during critical times. The method uses a combination of nature inspired algorithms, along with K-nearest neighbor for classification, achieved an impressive accuracy of 91.65%, outperforming earlier methods examined in the article, whereas in another paper, Jain et al. [10] used nature inspired algorithm. Sneha et al. [11], innovated a method for finding Diabetes Mellitus at an early stage through prognostic analysis by selecting relevant attributes. Results indicate that the Random Forest and decision tree classifier exhibited the uppermost specificity i.e. 98.20% and 98.00%, above all, making them suitable aimed at analyzing diabetes dataset. The model which achieved the highest accuracy 82.30% was Naïve Bayesian. The research also emphasizes the generalization of optimum feature selection as of the dataset to enhance the accuracy of classification model. In another approach, Dinh. et al. [8] employ machine learning to forecast cardiovascular diseases and diabetes using a data-driven strategy. Extreme gradient boosting is used in this study as a comparison to the LR, SVM, RF, and weighted ensemble models. The area under the ROC curve showed 95.7% accuracy. Nguyen et al. [12], used an ensemble classification model to predict type II diabetes. According to the AUC, the model's accuracy was 82.2%. Rajalakshmi et al. [13], developed a novel method for detecting diabetes using a smartphone. In this study, picture data was considered for diagnosis and future recommendations. Chen et al. [14] introduced a 5th Generation smart personalized diabetes prediction system for analyses and using 5G testbeds, smart phone, smart clothing, big data and clouds, the information and dataset was collected from a hospital in China.

Another perspective is, it can be difficult to extract useful data from electronic medical records and wireless sensor data for patients, to do interdisciplinary disease prediction. Different techniques are developed for the extraction of primary information from textual data related to healthcare to create a dataset for Type-2 diabetes prediction and have been presented by Bernal et al. [15]. To make the data rich and accurate, every wireless communication device data was preprocessed to diminish noise. Machine learning techniques will now be simple to use for predicting multimodal diabetic illness. Making a robust healthcare dataset for diabetes requires transforming

information from textual relevant data and combining it with already processed wireless sensor data [16] [17] [18]. The complete set of restrictions for previously used methods by other authors is shown below in Table 1

**Table 1. Relative work done by other authors**

| S. No | Author                   | Methods                                                                          | Novelty                                                                          | Outcomes                                                                       | Dataset                                           |
|-------|--------------------------|----------------------------------------------------------------------------------|----------------------------------------------------------------------------------|--------------------------------------------------------------------------------|---------------------------------------------------|
| 1     | Li et al. 2023 [6]       | Harmony search, Genetic, Particle swarm, KNN                                     | Use of AI enhancing tool for early diabetes prediction.                          | Early diabetes prediction using predictive analysis.                           | Diabetes dataset UCI machine repository           |
| 2     | Nguyen et al. 2019 [12]  | DL model with feed forward neural network                                        | Generalized linear model with DL features of feed forward neural network         | Improved performance of medical care using SMOTE                               | Electronic Health Record dataset of 9948 patients |
| 3     | Chen et al. 2018 [14]    | The 5G-Smart Diabetes testbed takes the guesswork out of managing your condition | 5G smart diabetes system for personalized analysis and treatment recommendations | Improved diabetes prediction with networking and intelligence                  | Hubei Province, China                             |
| 4     | Kaur et al. 2020 [18]    | R-data manipulation tool                                                         | Five ML models were used using R- tool for the accurate results                  | Develop and detect trends and patterns for the accurate prediction of diabetes | PIMA diabetes dataset                             |
| 5     | Yuan et al. 2019 [19]    | K-means and collaborative filtering (CF)                                         | An approach for personalized dietary needs of user                               | Improving the accuracy of maintaining nutritional balance in once diet         | Recorded dataset of 40 students                   |
| 6     | Sookrah et al. 2019 [20] | Content based filtering, Weka                                                    | Recommend personalized diet plan to hyper sensitive patients                     | Effective diet plan for hyper sensitive patients                               | USDA Food composition database                    |
| 7     | Eswari et al. 2015 [21]  | Hadoop/ Map-reduce                                                               | Predictive analysis of diabetes type to reduce the long term complications       | Analysis of diabetes type prevalent                                            | Electronic Health Record dataset                  |

|    |                            |                              | associated with it                                                                |                                                              |                              |
|----|----------------------------|------------------------------|-----------------------------------------------------------------------------------|--------------------------------------------------------------|------------------------------|
| 8  | Sarwar et al. 2018 [22]    | Ensemble model, MATLAB, Weka | Designed a framework to help and support doctors in biopsy                        | Support medical care by intelligent framework, timely biopsy | Survey dataset of 400 people |
| 9  | Annamalai et al. 2022 [23] | Salp swarm algorithm         | Hyper parameters selection using SSA for improving overall prediction performance | Identified severity of diabetes                              | PIMA diabetes dataset        |
| 10 | Ahlawat P. 2021 [24]       | KNN classifier               | Assist medical experts by providing a reliable and instant recommendation         | Solved imbalance problem, outliers and missing values        | PIMA diabetes dataset        |

Different types of mining techniques for data related to machine learning models, or blend of these techniques, have been used in various prediction models. A novel method for the diabetes data analysis system via Hadoop and Map Reduce was implemented in another paper [22]. The technique was used to forecast the type of diabetes and the likelihood of developing a serious illness. The technology is affordable for any healthcare company and is based on Hadoop [25][26]. One of the authors proposed a graph mining approach based on Random walk-based by Li et al. [27] that can also be used for staying healthy and can meet the user needs. In this study, Food Data Central (FDC), recipe internet site, and technical literature is collected and an assorted graph integrating all i.e. Food-Nutrition-Recipe Graph (FNRG) is constructed and proved beneficial for the people who suffered from chronic diabetes [28]. Chen et al. [29] projected a technique NutRec, the model gives healthily ingredient and its quantity for nutritious and healthy food recipe. The dataset used in this paper is from *yummly* and *allrecipe* and provided better results. Lopes-Nores et al. [30] proposed the solution to the problem which arises when traditional ways of recommendation apply to retrieve information from electronic health records. So, the author introduced a novel filtering strategy which differentiates user and items properties separately and the experiment comes out to more beneficial for the people with specific health concern. This novel approach solves the problem of sparsity, latency and inappropriate treatment given to people with different needs and interests. Various nature inspired algorithms are also used by the author for recipe recommendation. Jain and Singhal [31] created a system, which accurately recognized diabetes mellitus by applying nature-influenced algorithms to a specific dataset. Various classification algorithms are used for diabetes detection, and their accuracy is enhanced by hyper-parameters tuning using the Hybrid Bat metaheuristic algorithm. The hybrid ensemble classifier, with Smote and Bat Optimization Algorithm, outperforms others with a maximum accuracy of 98%. The primary focus of the author is on diabetes prediction, followed by an exploration of dietary recommendations.

Chen et al. [32] introduced the Food-Recommendation system which carries knowledge-aware attention graph convolutional neural network to capture the semantic relationships between users and recipes within the collaborative knowledge graph. By integrating the outcomes of these two learning tasks, FKGM is able to infer user preferences and health requirements. This work uses an Allrecipes dataset, and the achieved accuracy increased by 97.56%. A high error rate could not be achieved with this method. Mckensy-Sambola et al. [33] offered an efficient method called Ontology of Dietary Recommendations (ODR). The main goal of this method was to solve the problem of poor nutrition education. This research created a recommendation engine that can analyze users' levels of overweight and obesity and prescribe nutritional modifications to help them lose weight. This work uses a real-time dataset, and the achieved accuracy increased by 87%. The limitation of this method was to determine the complexity of the model.

Iwendi et al. [34] designed an efficient method called deep learning technique. This study leverages a health-based medical dataset that uses information about a patient's medical conditions as well as attributes. To determine which foods should be provided to which patients. This study used medical data collection for diabetics, and the accuracy improved by 97.74%. The drawback of this strategy was that the estimate step had a tiny sample size, and main meals were not considered. Table 2 shows the existing works related to the Diet recommendation.

**Table 2: Review related to the Diet Recommendation**

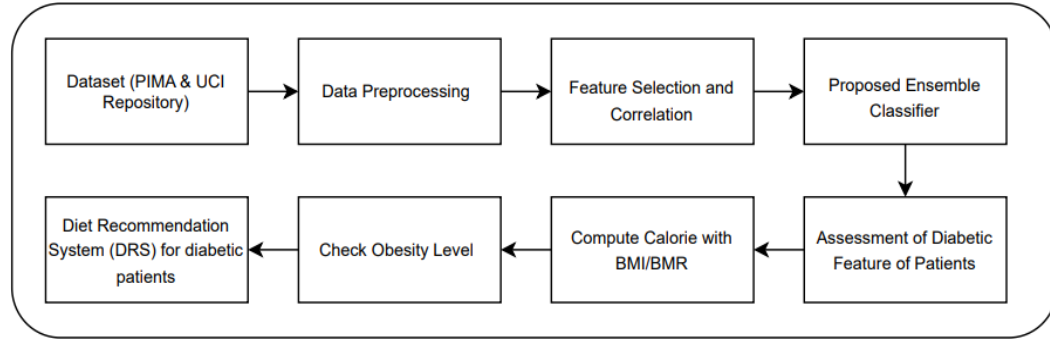
| Author details              | Methodology                                 | Repository                                                 | Challenge                               | Performance     |
|-----------------------------|---------------------------------------------|------------------------------------------------------------|-----------------------------------------|-----------------|
| Rostami et. al A [35]       | Deep Learning and Graph Clustering          | Publicly available food datasets                           | Low recall values                       | Accuracy-91%    |
| Shandilya et al. [36]       | MATURE-Health                               | Publicly available sodium, potassium, and BUN food dataset | Long training hours to train the models | Accuracy-99.53% |
| Rostami et. al B [37]       | deep learning-based image clustering method | Crawled dataset                                            | High time consumption.                  | Precision-7.35% |
| Chen et al. [32]            | FKGM                                        | Allrecipes dataset                                         | High error rate.                        | Accuracy-97.56% |
| Mckensy-Sambola et al. [33] | ODR                                         | Real-time data set                                         | Model complexity                        | Accuracy-87%    |

### 3. MATERIALS AND METHODS

In earlier times detection of diabetes was not easier for any medical expert and physician to do manually and accurately. To map dataset to specific categories various classifiers were used. The body mass index (BMI) of a person is one of the key indicators of whether they will develop diabetes. BMI does not measure body fat directly, but it comes out to be the strength of the paper as it shows health status of a diseased person who is diabetic whether normal weight, underweight or obese/overweight but does not account for reflection in their total body mass/fat and muscle conformation. Based on BMI, the author determines the basal metabolic rate (BMR), for the daily calorie intake needed by a diabetic person in accordance with that person's health category. Since a person's calorie intake needs to alter to maintain health, the BMI does not represent these variations.

The diabetes Mendeley dataset, PIMA, used the preprocessing techniques, and key features were selected for machine learning models. A feature selection approach identified important features, improving the performance of various classifiers such as the Ensemble Voting Classifier, Artificial Neural Network, AdaBoost, Random Forest, Support Vector Machine, Gradient Boost, and deep learning models like Tabnet and XgBoost. These models were then utilized to identify diabetes patients within the dataset. The framework, based on Sequential Feature Selection (SFS), was assessed using multiple performance metrics. The methodology for diabetes prediction is depicted in Figure 1.

To ensure model accuracy and reduce overfitting, 10-fold cross-validation was used, dividing the data into 10 equal parts, with each part being used once for validation and the remaining parts for training.



**Figure 1. Healthy Diet Recommendation workflow**

### 3.1 Proposed Ensemble Voting Classifier

An Ensemble classifier that predicts and classifies the data based on most votes after being trained on an ensemble of various models. In this study hybridization is designed with Random Forest Classifier [42], Support Vector Machine [43], AdaBoost Classifier [44], Gradient Boost [45] and AdaBoost Classifier [46], XGBoost, Tabnet, and ANN [47]. Voting classifiers are a composite of adopted classifiers [48] [49] [50] where, multiple classifiers are used for trained and their predictions are collectively fed into a voting mechanism to make a final prediction. A data set can be trained using a variety of techniques and an ensemble before predicting the result. According to both hard and soft voting, the majority decision on a forecast is made. The author's aim in this paper is to recommend food for diabetic patients based on their health status according to their health status whether underweight, overweight, and healthy. The well-being of diabetic patients is assessed based on their Body Mass Index (BMI); a parameter outlined in Table 3 of the training dataset. BMI is a measure calculated by dividing an individual's weight by their height. It's important to note that BMI doesn't capture variations in fat tissues and muscle volume. The 'w' mentioned is utilized to govern the weight of a specific person.

**Table 3: Health Status corresponding BMI Values**

| Individual's Health Status | BMI_Level Values |
|----------------------------|------------------|
| Under Weight Category      | 15 to 19.9       |
| Normal Weight Category     | 20 to 24.9       |
| Over Weight Category       | 25 to 29.9       |
| Obesity Level-I            | 30 to 34.9       |
| Obesity Level-II           | 35 to 39.9       |
| Obesity Level-III          | $\geq 40$        |

### 3.2 Datasets

The Mendeley PIMA Indian dataset includes important parameters crucial for the early detection of diabetes. It contains 768 records and features eight primary attributes that play a significant role in predicting diabetes. The outcome variable is binary, with 268 instances classified as diabetes cases and 500 as non-diabetes cases.

### 3.3 Algorithm

The step by step process is elaborated as follows:

Step 1. The reported datasets PIMA and UCL research data repository (Diet matrix) are separated into training and testing with 70% and 30% respectively.

Step 2. Opted Classifiers and Deep learning models, and initialized with input parameters as (C1, C2, ...Cn).

Step 3. Initial hyper parameters for classifiers are initialized as (C1, C2, ...Cn).

Step 4. Voting classifier is generated as ensemble classifier (Ci) based on weight selection.

Step 5. Hyper parameter tuning is performed using auto ML grid approaches.

Step 6. Final accuracy and Area Under Curve (AUC) is computed as performance measure.

The hybrid ensemble classifier yields 99.57% AUC and 98.85% accuracy which shows the proposed methodology has made significant contribution in this work. The methodology of the work is discussed in Algorithm 1.

| <b>Algorithm-1: Proposed Algorithm for Hybrid Ensemble Classifier</b> |                                                                                                                            |
|-----------------------------------------------------------------------|----------------------------------------------------------------------------------------------------------------------------|
| <b>Input:</b>                                                         | Dataset (D1, D2) contains Diabetic features and Diet matrix                                                                |
| <b>i.</b>                                                             | Initialize input parameters for classifiers (C1, C2, ...Cn)<br>Initialize hyper parameters for classifiers (C1, C2, ...Cn) |
| <b>ii.</b>                                                            | Generate Classifier Ci                                                                                                     |
| <b>iii.</b>                                                           | Obtain accuracy Vi for each Ci                                                                                             |
|                                                                       | Do                                                                                                                         |
| <b>iv.</b>                                                            | Compare performance of each Ci Assign weights for voting.                                                                  |
| <b>v.</b>                                                             | Hyper tuning of hybrid ensemble with Auto ML                                                                               |
| <b>vi.</b>                                                            | Compute the AUC                                                                                                            |
| <b>End</b>                                                            |                                                                                                                            |

$$BMI = \frac{(weight\ in\ kilograms)}{(height\ in\ meters)^2} \dots \dots \dots (1)$$

$$w = BMI_{normal} * height^2 \dots \dots \dots (2)$$

According to the provided equation, individuals can be classified as underweight, normal, or overweight. BMI serves as a measure of body weight status relative to fat (Equation 1 and Equation 2). It allows individuals to assess their excess body fat, associated health risks, and potential diseases related to carrying extra body fat. However, it's important to note that maintaining a desired weight and managing diabetes requires more than just considering BMI. Additionally, a noteworthy correlation exists between BMI and BMR (Basal Metabolic Rate). The sum of calories burned while resting is basal metabolic rate. BMR identified liable on a person's weight and height along with age. BMR provides the required count of calories needed to which a person must stick to regulate their diabetes. The Harris-Benedict equation shows that BMR is the same for both men and women as follows: The Harris-Benedict Equation gives your BMR to calculate your daily energy expenditure when at rest, after which an activity factor is added. Add the necessary activity component to your BMR as follows (Equation 3, 4, 5) [25]:

For Men:

$$BMR = 66.5 + (13.75 * w) + (5.003 * h) - (6.75 * age) \dots (3)$$

For Women:

$$BMR = 655.1 + (9.563 * w) + (1.850 * h) - (4.676 * age) \dots (4)$$

For diabetic individuals falling into specific BMI categories, their BMR is computed, and a recommended list of food items is provided for breakfast, lunch, and dinner 3 course meals grounded on their healthiness requirement. The daily calorie intake is distributed into 20%, 30%, and 50% for dinner, Lunch and breakfast respectively. For those with a BMI greater than 20, maintaining standard weight, and having diabetes, the recommended daily calorie intake is set between 1000-1200 calories.

### 3.4 Proposed Framework

The proposed framework demonstrates the prediction of diabetes and healthy diet recommendation model. The models discussed trained diverse attributes to predict either an individual is diabetic or not. If a person is identified as diabetic, a recommended-meal is provided. For food recommendations, a dataset from Kaggle is utilized. Diabetic individuals have their BMI checked to determine whether they are Underweight, Overweight, or Healthy. Subsequently, the daily calorie count required is calculated based on health status and BMI score. A curated list of food items is then suggested for breakfast, lunch, and dinner, aimed at controlling diabetes at an early stage. The workflow of this process is illustrated in Figure 2.

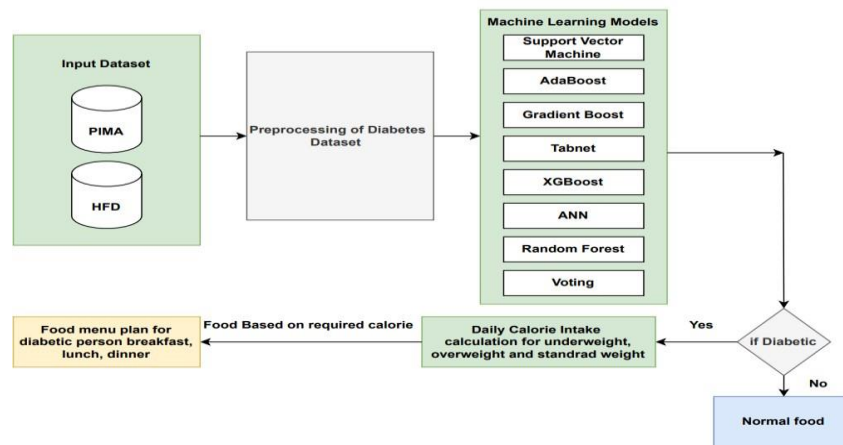


Figure 2. Proposed HDR framework Diabetes Prediction and Diet

## 4. EXPERIMENTAL SETUP

### 4.1 Simulation details

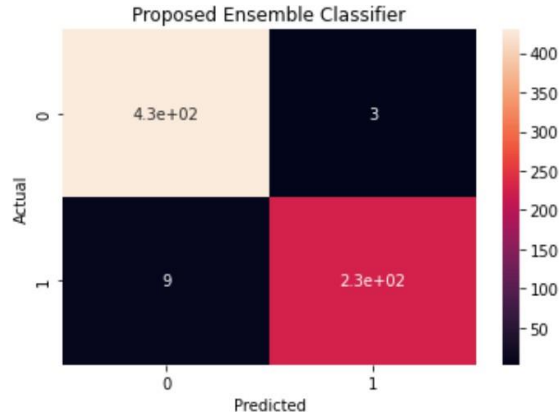
For implementation of the proposed framework, Sklearn of machine learning is used on Python 3. Sklearn provides a range of algorithms under supervised and unsupervised categories which includes all classifiers too. The details are shown in table 4.

### 4.2 Data Collection and Preprocessing

The classifiers performance measured using different parameters i.e., Confusion matrix, accuracy, precision, recall and F1-measures. The diabetes dataset is taken from *Mendeley PIMA India diabetes dataset*, and UCI food data set is taken from *Kaggle* that contains nutrient values like calories, fat, carbohydrates etc. [51] [52]. Diabetes dataset has specific parameters on which model is trained and tested, and prediction is done.

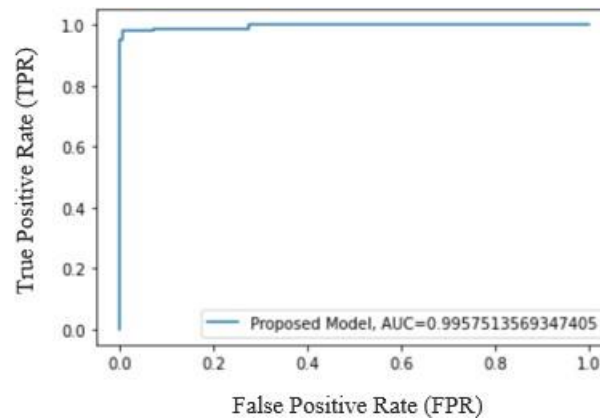
### 4.3 Performance Evaluation and Discussion

The model's overall effectiveness was assessed using heterogeneous indicators. The author has demonstrated the suggested learning model's overall predictive power using these criteria. The output of a classification problem, which can have more than two classes or more than two types of classes, is measured using a matrix [51].



**Figure 3. Proposed Ensemble Classifier**

The proposed model's confusion metric is shown in Figure 3 and the AUC curve is shown in Figure 4.



**Figure 4. Proposed Ensemble classifier Area under Curve (AUC)**

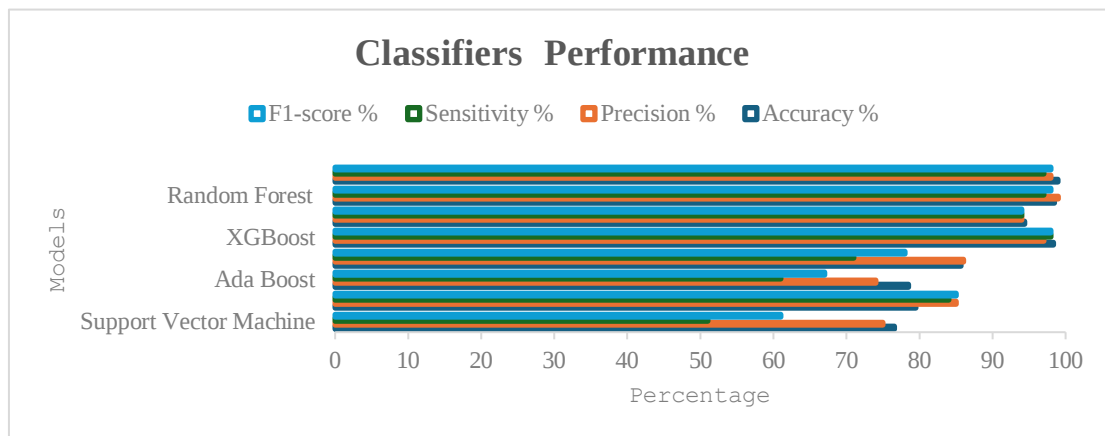
#### 4.4 Performance Metrics:

Table 4 depicts the performance of individual classifiers as Support vector machine, Tabnet, Ada-boost, Gradient boost, XGBoost, Artificial Neural Network, Random Forest compared with proposed model Ensemble Voting Classifier. The performance of the proposed hybrid ensemble classifier is examined through accuracy, precision, recall and F1-score. The proposed classifier improves accuracy with 98.85% and precision by 15% when compared other models.

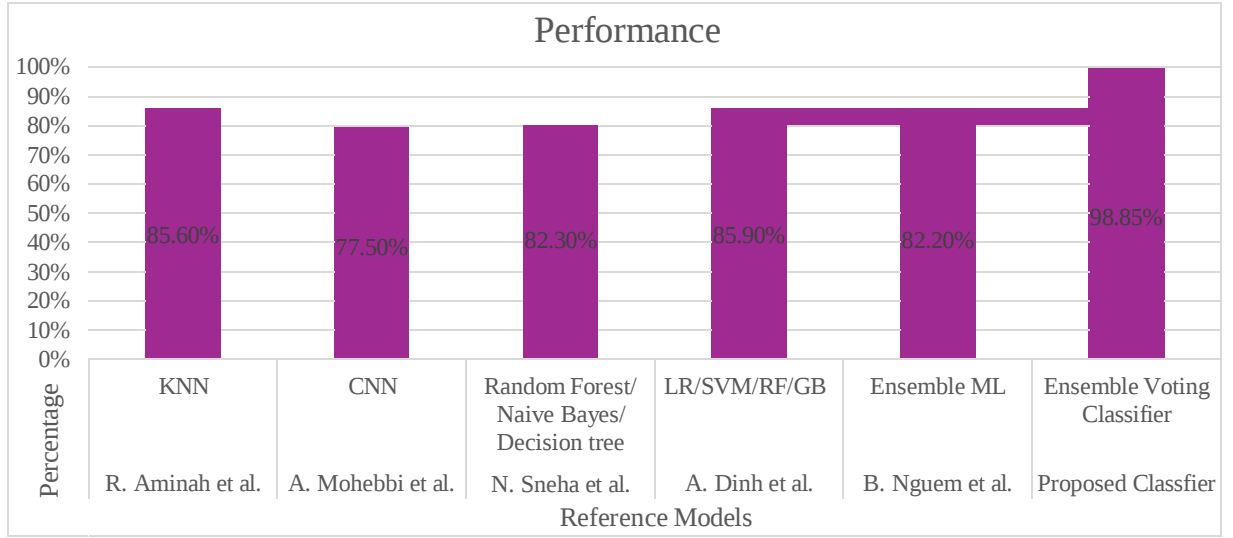
Several machine learning classifiers calculate the performance matrix and evaluate the correctness of the training data. Precision, recall, and F1-score and accuracy are used to evaluate one method against another. Methods used are Bagging, Random Forest, Support vector machines, Ada boost, Gradient boost, and Ensemble model discovered performance parameters. According to Table 4, the proposed classifier's accuracy is high when compared to other metrics.

**Table 4. Accuracy, Precision, Recall and F1-score values**

| Classifier's Performance   | Accuracy % | Precision % | Sensitivity % | F1-score % |
|----------------------------|------------|-------------|---------------|------------|
| Support Vector Machine     | 76.56      | 75          | 51            | 61         |
| Tabnet                     | 79.48      | 85          | 84            | 85         |
| Ada Boost                  | 78.50      | 74          | 61            | 67         |
| Gradient Boost             | 85.67      | 86          | 71            | 78         |
| XGBoost                    | 98.35      | 97          | 98            | 98         |
| Artificial Neural Network  | 94.40      | 94          | 94            | 94         |
| Random Forest              | 98.65      | 99          | 97            | 98         |
| Ensemble Voting Classifier | 98.85      | 98          | 97            | 98         |



**Figure 5: Performance of classifiers against individual models**



**Figure 6. Performance comparison with benchmark models**

A classifier's capacity to discriminate between classes is seen in Figure 5. The model exhibits enhanced performance in distinguishing between positive and negative classes as the area under the curve increases.

The comparison between the models utilized by other authors and the suggested model is shown in the figure 6. The illustration above demonstrates how well the suggested model worked with the ensemble model. The model proposed by the author has obtained maximum accuracy in comparison to previous studies. According to the method of prediction described above, the model's accuracy is 98.95% which is higher than that of other studies.

## 5. OPTIMIZED DIET RECOMMENDATION SYSTEM

This study proposes a nutrition recommendation system based on the triple attention method consisting of three attention modules. Initially, maintain half of the characteristics with higher weights by utilizing the traditional attention technique. The second approach, called average pooling attention, relies heavily on an average pooling function to hold on to neighborhood specifics. The third attention mechanism, max-pooling, is primarily employed to zero in on locally-relevant features that carry a heavier attention mechanism weight. Here is the formula for determining the sum of the three attention modules (*Equation 6, Equation 7, Equation 8*).

$$x = \sum_{b=1}^s \frac{\exp(z_{ab})}{\sum_{r=1}^s \exp(z_{ar})} n_{ab} \quad (6)$$

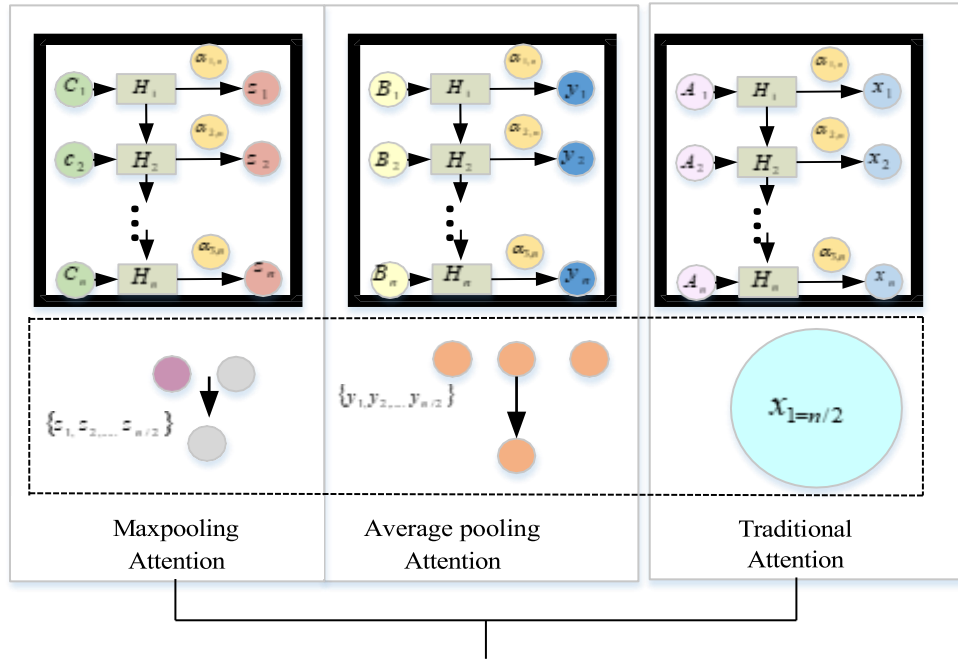
$$M = \text{AveragePooling} \left\{ \sum_{b=1}^s \frac{\exp(z_{ab})}{\sum_{r=1}^s \exp(z_{ar})} n_{ab} \right\} \quad (7)$$

$$l = \text{MaxPooling} \left\{ \sum_{b=1}^s \frac{\exp(z_{ab})}{\sum_{r=1}^s \exp(z_{ar})} n_{ab} \right\} \quad (8)$$

The triple attention strategy adds additional weights to each property of the original hyperspectral feature, ensuring that duplicate features cannot be removed in the future. Figure 7 shows the exact implementation approach for triple attention. The joint output equation for the triple-attention mechanism is as follows (*Equation 9*):

$$T_{(x,m,l)} = \text{concatenate}(x \oplus m \oplus l) \quad (9)$$

Here,  $x$  is denoted as the output of the original attention mechanism feature,  $m$  is denoted as the output of the maximum pooled attention mechanism feature,  $l$  is denoted as the average polled attention mechanism feature output,  $T_{(x,m,l)}$  is denoted as the joint output of triple attention feature, and the symbol  $\oplus$  is denoted as feature fusion algorithm operation. Fusion features are made by mixing different attention methods, and they can add more deep features to the residual dense network. To solve some of the flaws observed, the triple attention strategy was applied, resulting in a very successful technique that merged GRU and Bi-LSTM, as shown in Figure 7.



## Triple Attention Mechanism

**Figure 7: Triple attention mechanism structure**

Update and reset gates are accessible on the GRU [35]. The former gate is used by the GRU algorithm to determine whether to transmit past data. However, the model uses the later gate to decide how much previous information to forget, as mentioned in Equation 10 to Equation 13.

$$h_d = \{\text{sigmoid} * (G_h k_d + G_h S_{d-1} + V_h)\} \quad (10)$$

$$n_d = \{\text{sigmoid} * (G_h k_d + G_h S_{d-1} + V_n)\} \quad (11)$$

$$c'_d = \{\text{tanh} * (G_h k_d + G_h d \oplus S_{d-1})\} \quad (12)$$

$$c_d = \{H_d \oplus S_{d-1} + (1 - h_d) \oplus c'_d\} \quad (13)$$

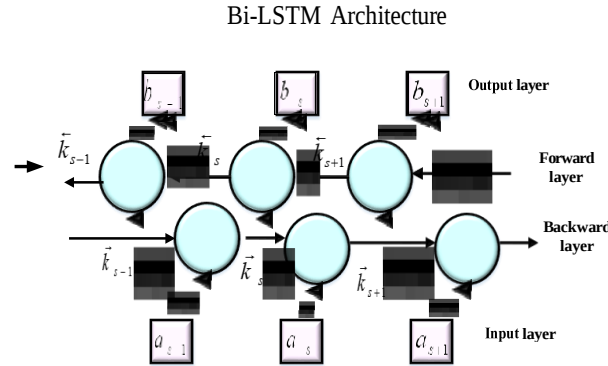
Here, the update gate is denoted as  $h_d$ , the reset gate is denoted as  $n_d$ , the candidate hidden state is denoted as  $c'_d$ , the new hidden state is denoted as  $c_d$ , the weight is denoted as  $G_h$ , the hidden state at the previous time step is denoted as  $S_{d-1}$ , the bias term is denoted as  $V_h$ , the hyperbolic activation is denoted as  $\text{tanh}$ , the Hadamard product is denoted as  $\oplus$ . To predict the current data, bidirectional models often have the unique ability to learn information from both past and subsequent (future) data. The key parameters and values considered for the proposed model are described in Table 5.

**Table 5: Input parameter values**

| Hyper-parameters                                       | Values     |
|--------------------------------------------------------|------------|
| Bidirectional GRU layer                                | 32         |
| First, second, third, and fourth convolutional layer   | 32,64,64,3 |
| Number of convolutional layer                          | 4          |
| Maximum pooling layer                                  | 2          |
| Activation function of first three convolutional layer | RELU       |

|                                    |                             |
|------------------------------------|-----------------------------|
| Activation function of final layer | Sigmoid                     |
| Optimizer                          | Artificial Rabbit Algorithm |
| Learning rate                      | 0.0002                      |
| Loss function                      | Mean absolute error         |
| Batch size                         | 16                          |
| Up sampling layer                  | 2                           |
| Number of epochs                   | 100                         |
| Dropout                            | 0.5                         |
| Epochs                             | 100                         |
| Verbose                            | 0                           |

The basic idea behind the Bi-LSTM models is to control cell states using three gates: input, forget, and output, as detailed in Figure 8.



**Figure 8: Bi-LSTM structure**

Through evaluating the input and hidden state values, the forget gate decides whether to keep or forget the data from the previous state, and its output value could be a 0 or 1. To check the cell state, the input gate determines how much information from the input text ( $\vec{a}_t$ ) and  $ht_1$  should pass, and its output could be a 0 or 1. The generated cell state is represented by the value of  $\vec{c}_t$  as a result of mathematical operations on  $\vec{c}_{t-1}$  and  $\vec{c}_t$ , as mentioned in Equation 14 to Equation 18.

$$g_t = \text{sigmoid}(X_{fy}z_t + X_{fg}j_{t-1} + v_f) \quad (14)$$

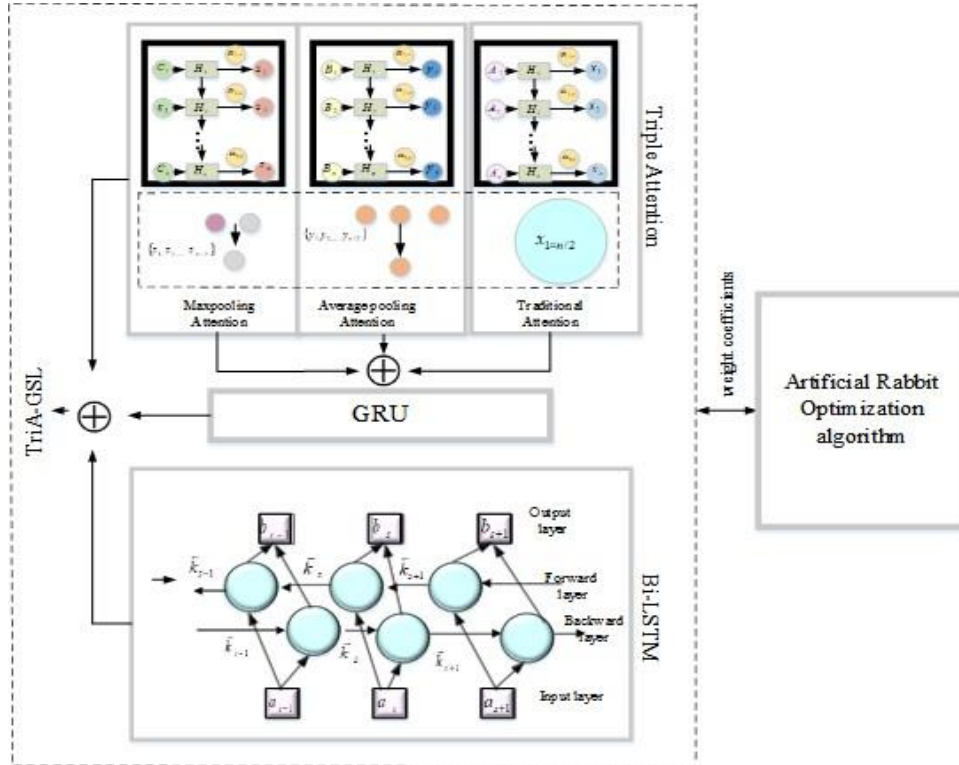
$$j_t = \text{sigmoid}(X_{fy}z_t + X_{fg}j_{t-1} + v_i) \quad (15)$$

$$a_t = a_{t-1} \odot g_t + j_t \odot \tanh(X_{bx}y_t + X_{bh}j_{t-1} + v_c) \quad (16)$$

$$p_t = \text{sigmoid}(X_{po}y_t + X_{ph}j_{t-1} + v_p) \quad (17)$$

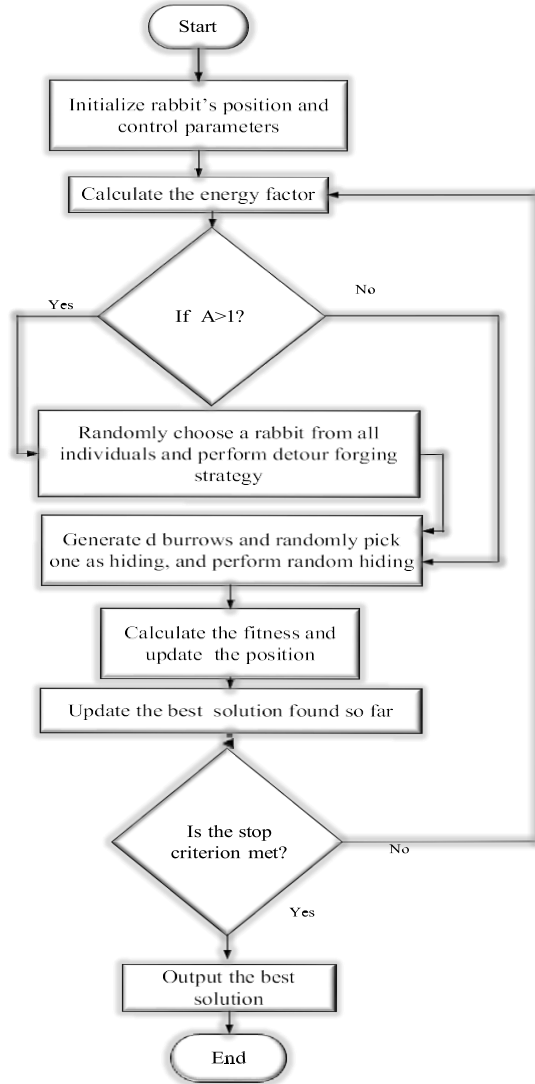
$$v_t = p_t \odot \tanh(a_t) \quad (18)$$

The output gate ( $p_t$ ), whose value may be 0 or 1, regulates the information flow from the active cell state to the hidden state.  $\vec{c}_t$  and the superscripts  $f$  and  $i$  show the input vector dimension and the datasets or vocabulary word count, respectively. The input vector, previously hidden state, and previous cell state are the inputs to LSTM at any given time, while the current hidden state and current cell state are the outputs. Multiplication of vectors using element by element is represented by the symbol  $\odot$ . The proposed TriA-GSL recommendation method is highly accurate to the diet recommendation shown in Figure 9.



**Figure 9: Proposed model recommendation structure**

The above TriA-GSL algorithm is for predicting specific problems, and calculating the weighting coefficients of each learning algorithm is a crucial point for the TriA-GSL methods. Therefore, an artificial rabbit optimization [37] (ARO) algorithm is developed to optimally obtain the weight coefficients of each classifier, which is better than the single learning algorithm. The flowchart of the ARO algorithm is shown in Figure 10.



**Figure 10:** Flow chart of Artificial Rabbit Optimization (ARO)

The ARO algorithm was inspired by the natural survival techniques adopted by rabbits. Detour foraging is the word used to describe rabbits' approach to finding food. The phrase "random hiding" describes how animals construct tunnels around their nests and occasionally hide inside of them to avoid predators and hunters. Rabbits can select between random hiding and detour foraging based on their energy level. If they have enough stamina, rabbits are able to leave far from their nests to seek food. Otherwise, they may slightly hide in neighboring burrows near their nests.

#### 5.1.1 Switch between Exploration and Exploitation

Rabbits have the ability to randomly hide or detour foraging. This is determined by the amount of energy a rabbit possesses; hence, an energy factor is provided (Equation 19).

$$X(y) = 4 \left(1 - \frac{y}{Y}\right) \lesssim \ln \frac{1}{d} \quad (19)$$

Here, selecting the random number between zero and one is denoted as the variable of .

#### 5.1.2 Detour Foraging

Predators are distracted from rabbit nests as the rabbits search for food. Based on the location of their neighbors, rabbits randomly search for food, as mentioned in Equation 20 to Equation 25.

$$\vec{c}_p(y+1) = \vec{o}_q(y) + D \times (\vec{o}_p(y) - \vec{o}_p(y)) + \text{round}(0.5 \times (0.05 + d_1)) \times f_1, \quad (20)$$

$$p, q = 1, \dots, f \text{ and } q \neq p \quad (21)$$

$$D = G \times h \quad (22)$$

$$G = (l - l^{\frac{d-1}{D}})^2 \times \sin(2\pi t_2) \quad (23)$$

$$A(u) = \begin{cases} 1 & \text{if } u = V(i) \\ 0 & \text{else} \end{cases} \quad m = 1, \dots, d \text{ and } i = 1, \dots, [d_3 \cdot t] \quad (24)$$

$$V = \text{randperm}(t) f_1 \sim D(0,1) \quad (25)$$

Here, the candidate position of  $p$ th rabbit at the time  $(y+1)$  is denoted as  $\vec{c}_p(y+1)$ , the  $i$ th rabbit position at the time  $r$  is denoted as  $\vec{o}_p(y)$ , the variable  $f$  is denoted as the rabbit population size, the maximum number of iterations is denoted as  $Y$ , the rabbit's movement of pace is denoted as  $i$ , the three random numbers between  $[0,1]$  is denoted as  $d_1$ ,  $d_2$  and  $d_3$ , the mapping vector denoted as  $A$ , and  $G$  is denoted as running operator that mimics the walk of a rabbit.

### 5.1.3 Random Hiding

When a predator approaches, it has the option of hiding in any of the holes that surround the rabbit's nest. In the following Equation 26 to Equation 31, each of the integers represents a distinct rabbit that created the burrows

$$\vec{z}_{p,q}(y) = \vec{o}_p(y) + J \times n \times \vec{o}_p(y), p = 1, \dots, f \text{ and } q = 1, \dots, t \quad (26)$$

$$J = \frac{D-d+1}{D} \times t_4 \quad (27)$$

$$n(u) = \begin{cases} 1 & \text{if } u = q \\ 0 & \text{else} \end{cases} \quad u = 1, \dots, t \quad (28)$$

$$\vec{c}_p(y+1) = \vec{o}_q(y) + D \times (t_4 \times \vec{z}_{p,q}(y) - \vec{o}_q(y)) p = 1, \dots, f \quad (29)$$

$$n_d(i) = \begin{cases} 1 & \text{if } i = [d_5 \times t] \\ 0 & \text{else} \end{cases} \quad i = 1, \dots, f \quad (30)$$

$$\vec{o}_p(y+1) = \begin{cases} \vec{o}_p(y) & \text{if } k(\vec{o}_p(y)) \leq k(\vec{z}_{p,q}(y+1)) \\ \vec{z}_{p,q}(y+1) & \text{if } k(\vec{o}_p(y)) > k(\vec{z}_{p,q}(y+1)) \end{cases} \quad (31)$$

Here, the hiding parameter is denoted as  $J$ ,  $\vec{z}_{pq}$  is denoted as  $j$ th burrow of the  $i$ th rabbit,  $\vec{z}_{pq}$  is denoted as the burrow for hiding for the  $i$ th rabbit, and  $t_4, t_5$  denoted as random numbers between  $(0,1)$ . The recommendation system uses a TriA-GSL model to predict accurate diet recommendations for diabetic patients and improve accuracy in less time.

## 6. Results

The various experimental investigations performed on the suggested and existing models are described in this section. The Mendeley PIMA dataset and the Diet Recommendation Dataset are two datasets that were utilized in this study to predict different classes using the diet recommendation system. For the two datasets, a number of performance measures are generated, including precision, recall, accuracy, F1-Score, sensitivity, root mean square error, Normalized Discounted Cumulative Gain, and mean absolute error. A deep learning technique is used to calculate accuracy and compare it to different current models, including VGG, ResNet, LSTM, Bi-LSTM, and the suggested model.

## 6.1 Dataset Description

The proposed model makes use of two datasets: the Mendeley PIMA dataset [39] and the UCI Diet Recommendation Dataset [40], which are described below.

### 6.1.1 UCI Dataset

Foods are contained in a dataset of, each containing one of the nine qualities listed as Glycemic Index, Grams, Protein, Carbohydrates, Dietary Fibre, Total Fat, Glycemic Load, Calories, and Fullness Factor. The food data is used to determine the glycemic load, calories and satiety factor. Based on this, the optimization algorithm determines the appropriate size for each portion.

### 6.1.2 Diet Recommendation Dataset

The dataset can be used to create an ML system that determines whether a particular type of food or range of packaged goods is suitable for a particular customer to maintain a nutritious diet for the patient. The foods included in the dataset are all necessary for human survival, and each has the following characteristics: meals, including breakfast, lunch and dinner; vegetarian or non-vegetarian alternatives; carbs, lipids, proteins, minerals, calcium, salt, and sodium; fibre, and sugars.

## 6.2 Performance Metrics

To evaluate the performance of a proposed categorization model, a range of performance measures are used. The following are the calculations for the various metrics:

### 6.2.1 Precision:

Precision measures the total number of samples that the classifier properly categorized or all of the true positives that were discovered during classification. The formulation is as follows (*Equation 32*):

$$P = \frac{ab}{ab+cb} \quad (32)$$

Here, ab is denotes the true positives, and cb is denotes the false positives.

### 6.2.2 Recall:

Recall is a measure that was created based on the true positives and true negatives found during categorization. It illustrates the distinction between information that was accurately and incorrectly categorized. An illustration of the formulation is shown in *Equation 33*:

$$R = \frac{ab}{ab+cd} \quad (33)$$

Here, cd is indicates the false negative values.

### 6.2.3 F-measure:

Based on the classification truth values, the F-measure is calculated using the harmonic mean of the accuracy and recall values. Since this number is more reliant on categorization accuracy, a high value for this indicator is desired. The computation of the f-measure may be written in *Equation 34*:

$$F1 = \frac{2*P*R}{P+R} \quad (34)$$

### 6.2.4 Accuracy:

Accuracy is the single most important measure of a recommendation's effectiveness. The formula for accuracy is mentioned in *Equation 35*:

$$A = \frac{ab+ad}{ab+ad+cd+cb} \quad (35)$$

6.2.5 Root mean square error (RMSE):

RMSE is a measure of error that reveals the overall misclassifications made by the model (*Equation 36*).

$$RMSE = \sqrt{\frac{\sum_{p=1}^m (x_p - \hat{x}_p)^2}{m}} \quad (36)$$

Here,  $x_p$  is denoted as prediction,  $\hat{x}_p$  is denoted as true value, and  $m$  is denoted as the total number of data points.

6.2.6 Mean absolute error (MAE):

The average variation between the actual and measured values is referred to as MAE, and the mathematical equation for this is given below in *Equation 37*:

$$MAE = \frac{\sum_{p=1}^m abs|x_p - \hat{x}_p|}{m} \quad (37)$$

6.2.7 Normalized Discounted Cumulative Gain (NDCG):

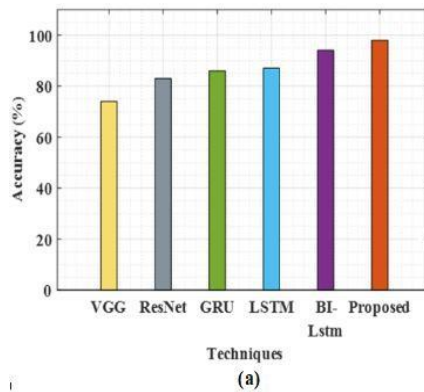
The NDCG assigns greater weight to hits that appear higher in the list of rankings. It is more likely that the suggestion list will perform well for relevant items if the NDCG score is high.

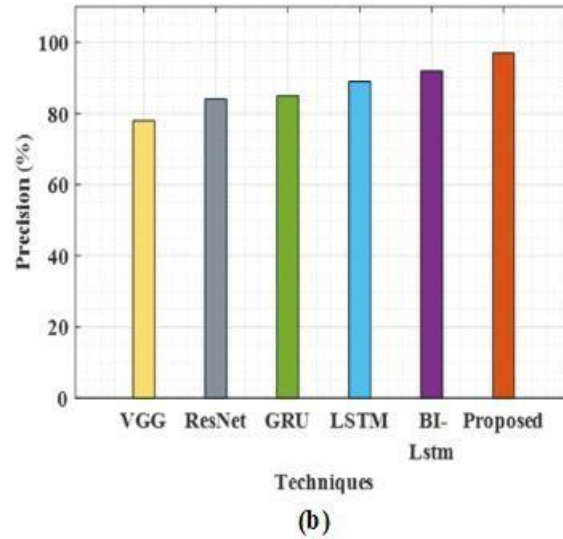
### 6.3 Performance analysis

To evaluate performance improvement, the suggested method has been contrasted with alternative recommendation models applied to the same dataset. The section below provides an illustration of the performance analysis of the suggested and current models over the UCI dataset and the Diet Recommendation Dataset.

6.3.1 Evaluation of UCI dataset

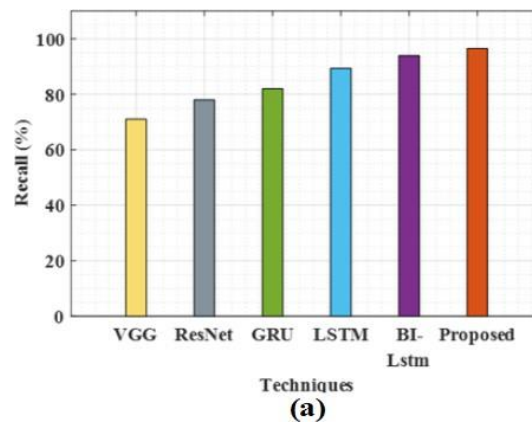
An investigation of a UCI dataset revealed that the proposed deep learning approach enhanced the overall performance of a recommendation stage, resulting in better and more precise identification of dietary suggestions for diabetic patients. Overall results suggest that the recommended strategy is better than other approaches for proposing a diabetic patient. Figures 11(a) and 11(b) compare the accuracy and precision of the suggested model with the existing model.

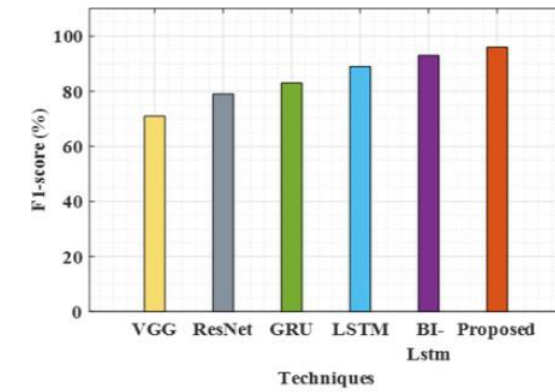




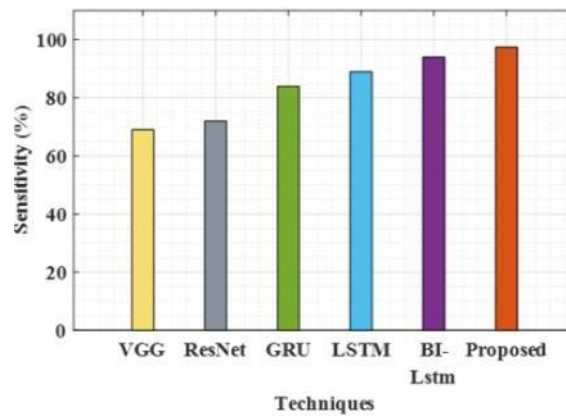
**Figure 11 (a) and 11 (b):** Analysis of accuracy and precision for both the existing and suggested models

The accuracy and precision compared to the suggested and existing models are shown in Figures 11(a) and 11(b). In comparison to the existing recommendation system, the suggested recommendation system has a 98.09% overall accuracy rate and a 97.65% precision rate. For the proposed recommendation model to work, the results of the performance described above must be extremely impressive. The existing model only accomplished a modest range due to limitations such as high computation times and low recommendation accuracy. The suggested recommendation model improves the drawbacks of the existing model, outperforming existing strategies and delivering outcomes with greater accuracy and precision. Figures 12(a), 12(b), and 12(c) compare the suggested model with the existing model.





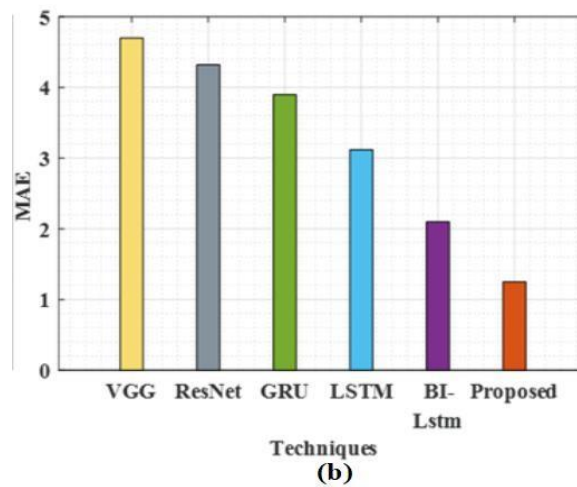
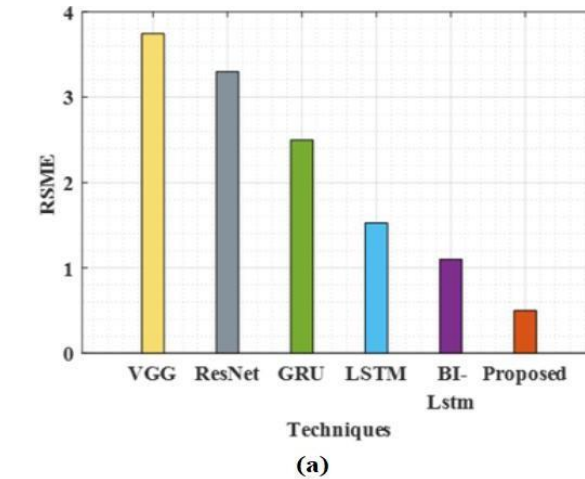
(b)



(c)

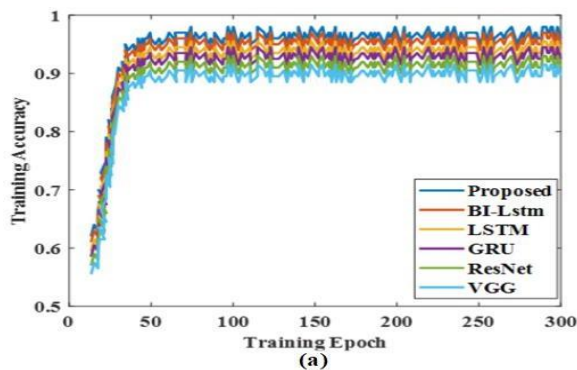
**Figure 12(a), 12(b), and 12(c):** Analysis of recall, f1-score, and sensitivity for both the existing and suggested models

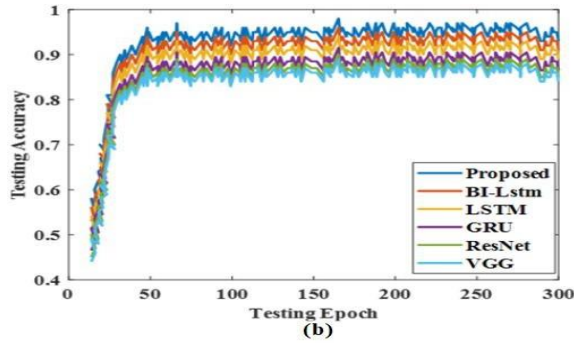
The suggested recommendation model was then compared against the performance of the existing recommendation model. The proposed model achieved high levels of recall, sensitivity, and F1-measure at 96.51%, 97.05%, and 96.8%. The outcomes of the performance given above must be highly varied for the suggested recommendation model to be effective. Due to limitations like limited recall values and model complexity, the current model had achieved limited coverage. The suggested recommendation model addresses the weakness of the existing model, outperforming existing approaches and producing superior results. Figure 13(a) and (b) shows the comparison of the suggested model's RMSE and MAE performance.



**Figure 13(a) and 13(b):** Analysis of the suggested TriA-GSL recommendation system error matrices

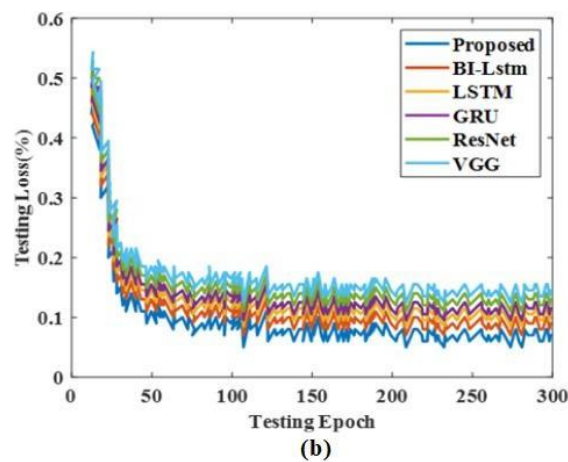
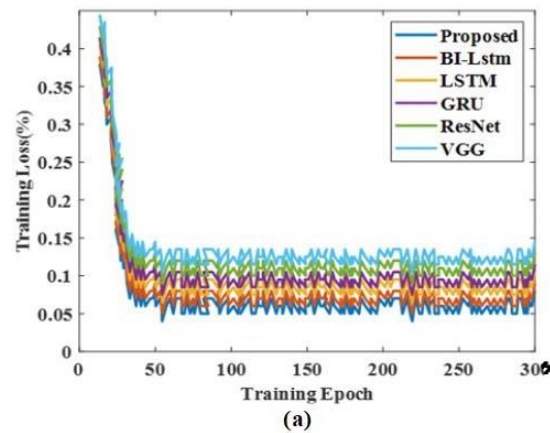
The suggested recommendation system can be analyzed through two error matrices, RMSE as 0.5 and MAE as 1.25, that outperformed the other compared suggested recommendation techniques. The RMSE and MAE of VGG are attained to be 3.745 and 4.7, ResNet as 3.3 and 4.32, GRU as 2.5 and 3.9, LSTM as 1.526 and 3.12 and Bi-LSTM as 1.1 and 2.1, respectively. So, the suggested recommendation system has obtained a low error performance compared to the existing models. The suggested recommendation technique training and testing accuracy are depicted in Figures 14(a) and 14(b).





**Figure 14(a) and 14(b):** Accuracy analysis for training and testing

The existing recommendation and suggested recommendation models are capable of obtaining values of 0.69 and 0.9870 during iterations of epoch values at 300 in the training set. Because the recommended technique used training data for several iterations to increase accuracy, it can achieve better accuracy than existing models. The tested set's existing and suggested models can provide values of 0.534 and 0.9489 for epoch values at 300 iterations. Because the suggested method tests data multiple times to enhance accuracy, it can achieve higher accuracy than existing models. The training and testing loss recommended method is depicted in Figures 15(a) and 15(b).

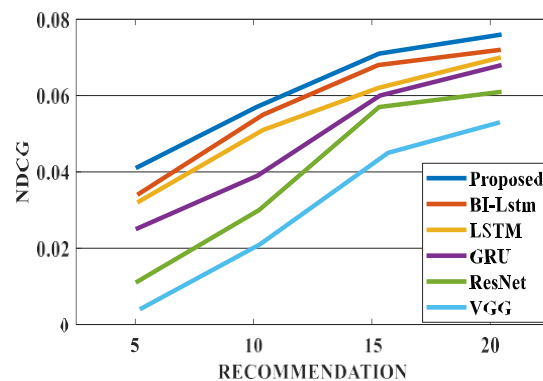


**Figure 15(a) and 15(b):** Loss analysis for training and testing

The existing recommendation and suggested recommendation models can obtain values of 0.456 and 0.05 during iterations of epoch values at 300 in the training set. Because the recommended technique trained data over Multiple iterations to improve the loss rate, it can achieve a lower training loss rate than existing models. In the

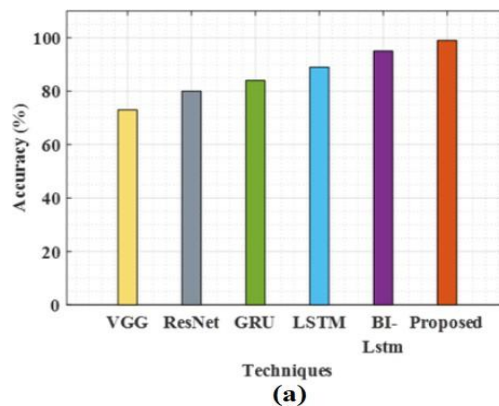
testing set, the existing and suggested models may provide values of 0.567 and 0.156 for epoch values at 300 iterations, respectively. Because the proposed method evaluates data for multiple iterations to lessen testing loss, it can therefore achieve lower testing loss than existing models. Figure 16 compares the NDCG of the recommended model with the existing model.

Figure 16 displays that NDCG measured and compared with the existing recommendation and suggested recommendation model. Compared to the NDCG of the existing recommendation system, the total NDCG of the suggested recommendation system is 0.076%. Due to the existing system's low NDCG value, the suggested recommendation system gave higher NDCG values during compression than the comparable approach



**Figure 16:** Analysis of the NDCG

### 6.2.3 Evaluation of Diet Recommendation Dataset



An investigation of the Diet Recommendation dataset revealed that the proposed deep learning technique enhanced the overall performance of a recommendation stage, resulting in better and more accurate identification of dietary suggestions for diabetes patients. Overall results suggest that the suggested strategy for advising a diabetic patient's diet is better than alternative approaches. Figure 17(a) and 17(b) compares the accuracy and precision of the recommended model with the present model.

Figure 17(a) and 17(b): Analysis of accuracy and precision for both the existing and suggested models

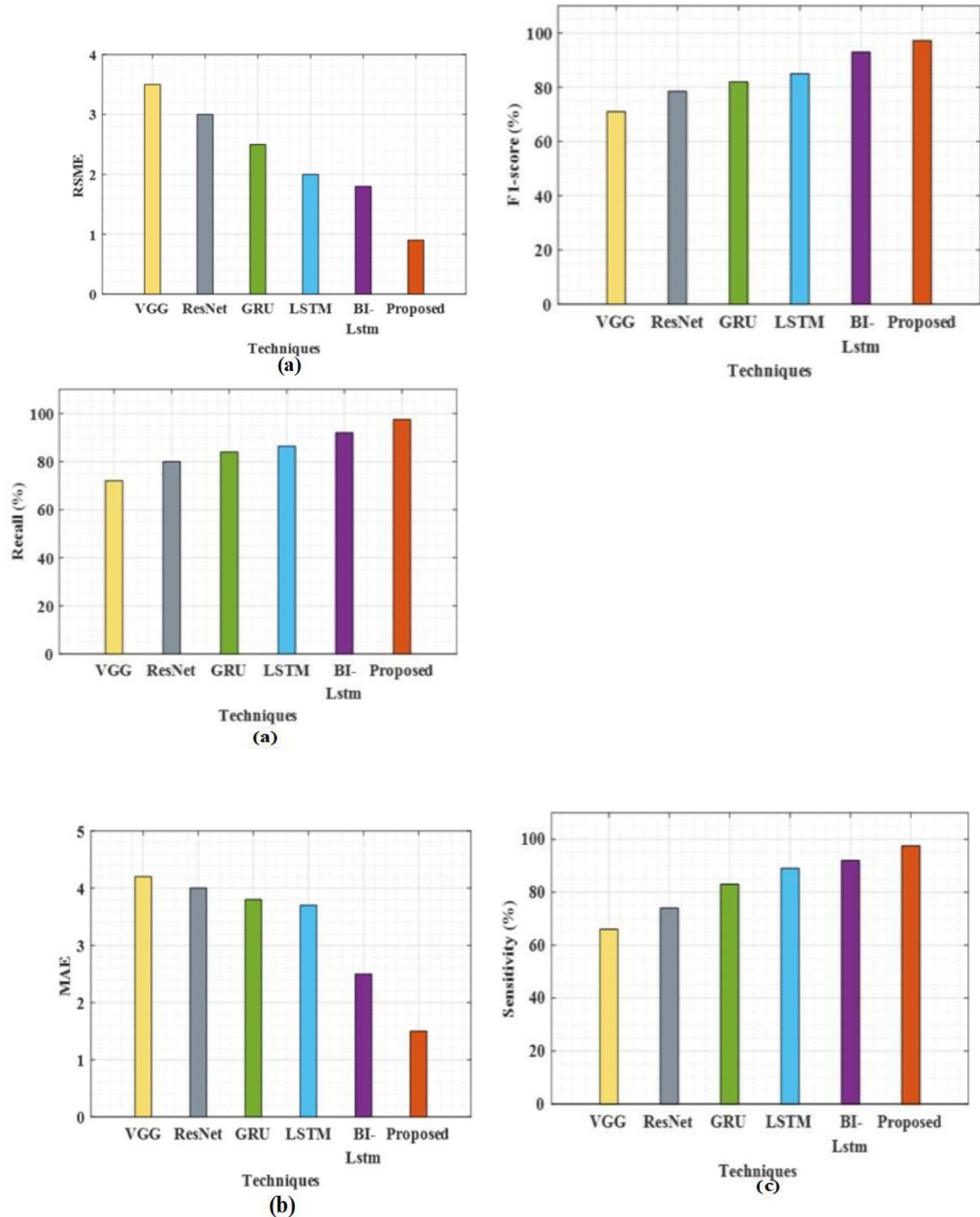


Figure 18(a), 18(b) and 18(c): Analysis of recall, f1-score, and sensitivity for both the existing and suggested models

## 6.4 Discussion

In comparison to the existing recommendation system, the suggested recommendation system has a 99.45% overall accuracy rate and a 98.05% precision rate. The existing recommendation model only accomplished a modest

range due to limitations such as high computation times and low recommendation accuracy. The suggested recommendation model improves the drawbacks of the existing recommendation model, outperforming existing recommendation strategies and delivering outcomes with greater accuracy and precision. Figures 18(a), 18(b), and 18(c) compare the recall, f1-score, and sensitivity of the suggested recommended model with the existing recommendation model. The effectiveness of the current recommendation model was then contrasted with the suggested recommendation model. The proposed model achieved high levels of recall, sensitivity, and F1-measure, including 97.65%, 98.09%, and 97.28%. The outcomes of the performance given above must be highly rare for the recommendation model to be effective. Due to limitations such as limited recall values and model complexity, the current model had limited coverage. The suggested recommendation model addresses the weakness of the existing model, outperforming existing approaches and producing superior results. Figure 19(a) and 19(b) shows the comparison of the suggested model's RMSE and MAE performance.

Figure 19(a) and 19(b): Analysis of the suggested recommendation system error matrices

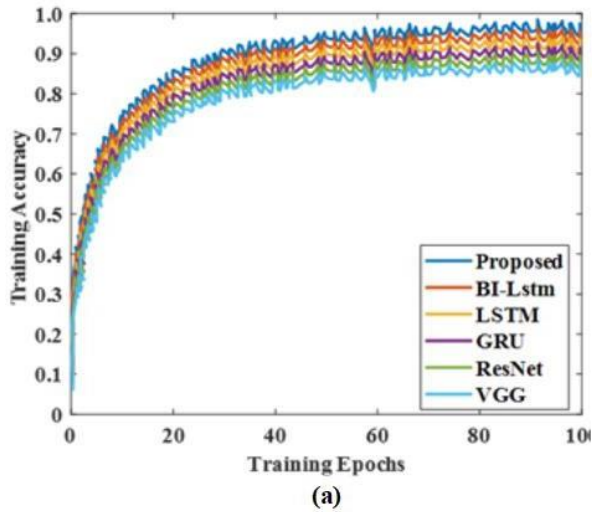


Figure 19(a) and 19(b), the suggested recommendation system can be analyzed through two error matrices: RMSE as 0.9 and MAE as 1.5 outperformed the other compared suggested recommendation techniques. The RMSE and MAE of VGG are attained to be 3.5 and 4.2, ResNet as 3 and 4, GRU as 2.5 and 3.8, LSTM as 2 and 3.7, and Bi-LSTM as 1.8 and 2.5, respectively. So, the suggested system has obtained a low error performance compared to the existing models. The suggested technique training and testing accuracy are depicted in Figures 20(a) and 20(b).

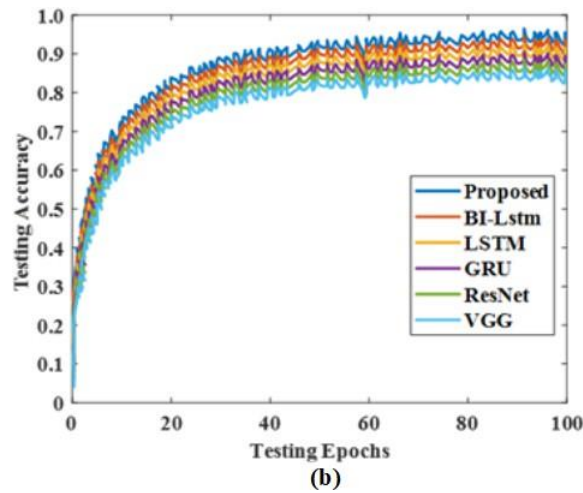
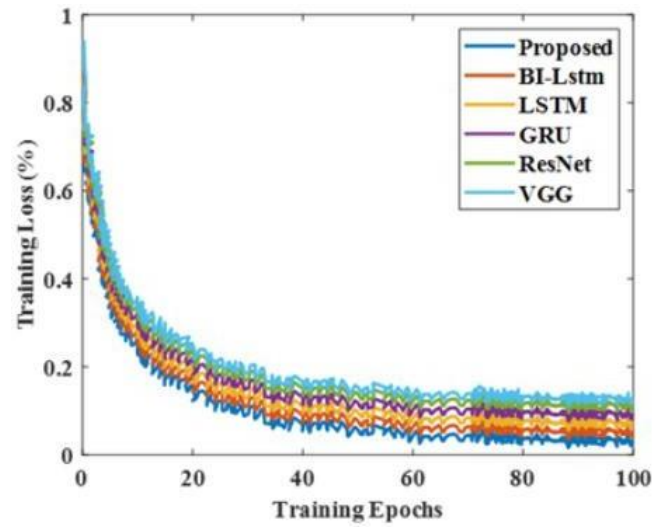
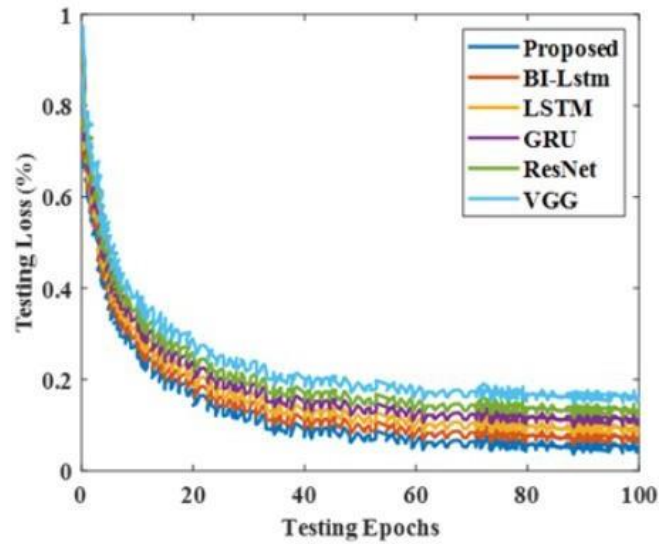


Figure 20(a) and 20(b): Accuracy analysis for training and testing

The existing and proposed models can produce values between 0.09 and 0.99 during iterations of epoch values at 300 in the training set. Because the suggested recommended technique used training data for several iterations to increase accuracy, it can achieve better accuracy than existing recommendation models. The tested set's existing recommendation and suggested recommendation models can provide values of 0.086 and 0.954 for epoch values at 300 iterations. During iterations of epoch values at 300 in the training set, the existing and proposed models can produce values ranging from 0.09 to 0.99. The suggested recommendation approach to training and testing loss are shown in Figures 21(a) and 21(b).



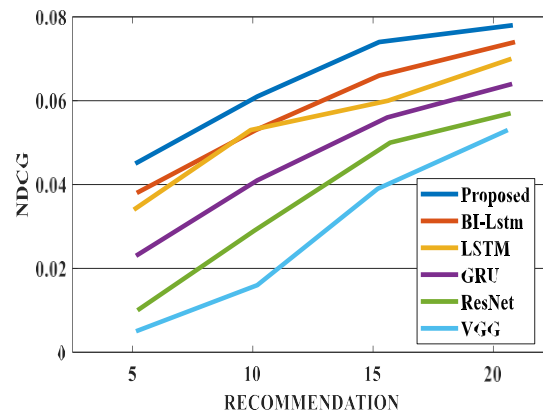
(a)



(b)

**Figure 21(a) and 21(b):** Loss analysis for training and testing

The current and suggested models can obtain values of 0.96 and 0.003 during iterations of epoch values at 300 in the training set. Because the suggested recommended technique trained data over multiple iterations to improve the loss rate, it can achieve a lower training loss rate than existing recommendation models. In the testing set, the existing and suggested models may provide values of 0.1 and 0.01 for epoch values at 300 iterations, respectively. As a result of testing data for many iterations to improve testing loss, the proposed recommendation approach can achieve lower testing loss than existing recommendation models. Figure 22 compares the NDCG of the suggested recommended model with the existing recommendation model.



**Figure 22:** Analysis of the NDCG

Figure 22 displays NDCG measured and compared with the existing recommendation and suggested recommendation model. Compared to the NDCG of the existing recommendation system, the total NDCG of the suggested recommendation system is 0.074%. Due to the existing recommendation system's low NDCG value, the suggested recommendation system gave higher NDCG values during compression than the comparable approach.

## 7. CONCLUSION

The suggested study demonstrates the early diabetes prediction and recommendation management of the diet based on the user's input values for early diabetes control. The Mendeley PIMA online repository and the UCI food data were used as input data set. The suggested hybrid ensemble classifier determines whether a patient is diabetic or not by looking at the characteristics of the disease. Model correctness is calculated using a variety of performance metrics. In addition to this, the study implements a TriA-GSL approach for the recommendation system in an effort to improve its efficacy and precision. This method establishes an efficient framework for recommending breakfast, lunch, and dinner for diabetic patients. This recommendation system is accomplished through several steps, including pre-processing, similarity calculation, and clustering and recommendation system. To recommend text, two publicly valuable datasets, such as Mendeley PIMA and the UCI Diet recommendation dataset, are used. A range of performance parameters are assessed and contrasted with pre-existing and proposed models for two distinct datasets. The values that can be obtained by the proposed model are as follows: 0.05 for training loss, 0.156 for testing loss, 0.9870 for training accuracy, and 0.9489 for testing accuracy. The proposed model parameter accuracy 98.09%, precision 97.65%, recall 96.51%, F1score 96.8%, NDCG 97.05%, RSME 0.5%, and MAE 1.25%, are also measured in the PIMA dataset. For a Diet recommendation Dataset, the proposed model can obtain values in training loss and testing loss of 0.003 and 0.01. The Diet recommendation dataset can obtain values of 0.99 to 0.954 for the proposed model in training and testing accuracy. The proposed model can obtain high accuracy compared to existing models. The suggested model's high time consumption and complexity are its drawbacks. Therefore, the nutritional qualities of each item can be added to the proposed work as extra information, and dietary suggestions based on each person's health status and diseases would be provided.

## LIST OF ABBREVIATIONS

CNN = Convolutional Neural Network

LR = Linear Regression

MLP = Multi-Layer Perceptron

NDCG = Normalized Discounted Cumulative Gain

AUC = Area under Curve

ML = Machine Learning

## **AUTHORS CONTRIBUTION**

I, [Anjali Jain], declare that this research paper titled " " is my own original work. All the ideas, concepts, and data presented herein are a result of my independent research efforts. I have duly acknowledged all the sources of information used in this paper through proper citation. Dr alka Singhal has proof read the paper

## **CONSENT FOR PUBLICATION**

Authors are giving all rights to the publishers to publish and disseminate the article as per the policy and agreement.

## **AVAILABILITY OF DATA AND MATERIALS**

The source of data and materials is cited in the manuscript. The reference is mentioned in the reference section. Also, the data will be available on demand via mail to the corresponding author.

## **FUNDING**

No funding agencies

## **Conflict of Interest**

The authors declare no conflict of interest, financial or otherwise.

## **Acknowledgements**

Authors acknowledge the role of the Jayee institute and KIET college in the research work.

## **REFERENCES**

1. Artzi, N. S., Shilo, S., Hadar, E., Rossman, H., Barbash-Hazan, S., Ben-Haroush, A., & Segal, E. (2020). Prediction of gestational diabetes based on nationwide electronic health records. *Nature medicine*, 26(1), 71-76.
2. Aminah, R., & Saputro, A. H. (2019, September). Diabetes prediction system based on iridology using machine learning. In 2019 6th International Conference on Information Technology, Computer and Electrical Engineering (ICITACEE) (pp. 1-6). IEEE.
3. Jain, A., & Singhal, A. (2022) Diet Recommendation using Predictive Learning Approaches, 3rd International Conference on Issues and Challenges in Intelligent Computing Techniques (ICICT), Ghaziabad, India, 2022, pp. 1-5.
4. Li, Y., Zhan, Z., Li, H., & Liu, C. (2021). Interest-aware influence diffusion model for social recommendation. *Journal of Intelligent Information Systems*, 1-15.
5. Mohebbi, A., Aradottir, T. B., Johansen, A. R., Bengtsson, H., Fraccaro, M., & Mørup, M. (2017, July). A deep learning approach to adherence detection for type 2 diabetics. In 2017 39th Annual International Conference of the IEEE Engineering in Medicine and Biology Society (EMBC) (pp. 2896-2899). IEEE.
6. Li, X., Zhang, J. & Safara, F. Improving the Accuracy of Diabetes Diagnosis Applications through a Hybrid Feature Selection Algorithm. *Neural Process Lett* 55, 153–169 (2023). <https://doi.org/10.1007/s11063-021-10491-0>.
7. Nallarasan, V., Anand, L., & Prabakaran, J. IMPROVED PREDICTIVE LEARNING APPROACHES FOR CUSTOMIZED DIET SUGGESTION FRAMEWORK IN MEDICINAL SERVICES. *Journal of Xi'an University of Architecture & Technology*, 2023.
8. Dinh, A., Miertschin, S., Young, A., & Mohanty, S. D. (2019). A data-driven approach to predicting diabetes and cardiovascular disease with machine learning. *BMC medical informatics and decision making*, 19(1), 1-15.
9. PIMA Dataset, accessed on 21 Jan 2025, <https://data.mendeley.com/datasets/7zcc8v6hvp/1>
10. Jain, A., Singhal A., (2024). Bio-inspired Approach for Early Diabetes Prediction and Diet Recommendation. *SN Computer Science*, 5(182).

11. Sneha, N., Gangil, T. Analysis of diabetes mellitus for early prediction using optimal features selection. *J Big Data* 6, 13 (2019). <https://doi.org/10.1186/s40537-019-0175-6>
12. Nguyen, B. P., Pham, H. N., Tran, H., Nghiem, N., Nguyen, Q. H., Do, T. T., ... & Simpson, C. R. (2019). Predicting the onset of type 2 diabetes using wide and deep learning with electronic health records. *Computer methods and programs in biomedicine*, 182, 105055.
13. Rajalakshmi, R., Subashini, R., Anjana, R. M., & Mohan, V. (2018). Automated diabetic retinopathy detection in smartphone-based fundus photography using artificial intelligence. *Eye*, 32(6), 1138-1144.
14. Chen, M., Yang, J., Zhou, J., Hao, Y., Zhang, J., & Youn, C. H. (2018). 5G-smart diabetes: Toward personalized diabetes diagnosis with healthcare big data clouds. *IEEE Communications Magazine*, 56(4), 16-23.
15. Bernal, E. A., Yang, X., Li, Q., Kumar, J., Madhvanath, S., Ramesh, P., & Bala, R. (2017). Deep temporal multimodal fusion for medical procedure monitoring using wearable sensors. *IEEE Transactions on Multimedia*, 20(1), 107-118.
16. Muzammal, M., Talat, R., Sodhro, A. H., & Pirbhulal, S. (2020). A multi-sensor data fusion enabled ensemble approach for medical data from body sensor networks. *Information Fusion*, 53, 155-164.
17. Nweke, H. F., Teh, Y. W., Alo, U. R., & Mujtaba, G. (2018, May). Analysis of multi-sensor fusion for mobile and wearable sensor based human activity recognition. In *Proceedings of the international conference on data processing and applications* (pp. 22-26).
18. Kaur, H., & Kumari, V. (2020). Predictive modelling and analytics for diabetes using a machine learning approach. *Applied computing and informatics*.
19. Yuan, Z., & Luo, F. (2019, June). Personalized Diet Recommendation Based on K-means and Collaborative Filtering Algorithm. In *Journal of Physics: Conference Series* (Vol. 1213, No. 3, p. 032013). IOP Publishing.
20. Sookrah, R., Dhowtal, J. D., & Nagowah, S. D. (2019, July). A DASH diet recommendation system for hypertensive patients using machine learning. In *2019 7th International Conference on Information and Communication Technology (ICoICT)* (pp. 1-6). IEEE.
21. Eswari, T., Sampath, P., & Lavanya, S. J. P. C. S. (2015). Predictive methodology for diabetic data analysis in big data. *Procedia Computer Science*, 50, 203-208.
22. Sarwar, A., Ali, M., Manhas, J., & Sharma, V. (2020). Diagnosis of diabetes type-II using hybrid machine learning based ensemble model. *International Journal of Information Technology*, 12(2), 419-428.
23. Annamalai, R., & Nedunchelian, R. (2022). Design of optimal bidirectional long short-term memory based predictive analysis and severity estimation model for diabetes mellitus. *International Journal of Information Technology*, 1-9
24. Ahlawat, P. (2021). DCPM: An effective and robust approach for diabetes classification and prediction. *International Journal of Information Technology*, 13(3), 1079-1088.
25. Osadchiy, T., Poliakov, I., Olivier, P., Rowland, M., & Foster, E. (2019). Recommender system based on pairwise association rules. *Expert Systems with Applications*, 115, 535-542.
26. Rehman, F., Khalid, O., Bilal, K., & Madani, S. A. (2017). Diet-right: A smart food recommendation system. *KSII Transactions on Internet and Information Systems (TIIS)*, 11(6), 2910-2925.
27. Li, K., Jiang, Z., Wang, H., & Liu, X. (2020). Healthy diet recommendation via Food-Nutrition-Recipe Graph mining. *Proceedings of the Association for Information Science and Technology*, 57(1), e319.
28. Namgung, K., Kim, T. H., & Hong, Y. S. (2019). Menu recommendation system using smart plates for well-balanced diet habits of young children. *Wireless Communications and Mobile Computing*, 2019.
29. Chen, M., Jia, X., Gorbonos, E., Hoang, C. T., Yu, X., & Liu, Y. (2020). Eating healthier: Exploring nutrition information for healthier recipe recommendation. *Information Processing & Management*, 57(6), 102051.
30. López-Nores, M., Blanco-Fernández, Y., Pazos-Arias, J. J., & Gil-Solla, A. (2012). Property-based collaborative filtering for health-aware recommender systems. *Expert Systems with Applications*, 39(8), 7451-7457.
31. Jain, A., & Singhal, A. (2023, November). Utilizing Metaheuristic Machine Learning Techniques for Early Diabetes Detection. In *2023 Second International Conference on Informatics (ICI)* (pp. 1-6). IEEE.
32. Y. Chen, Y. Guo, Q. Fan, Q. Zhang, Y. Dong, Health-Aware Food Recommendation Based on Knowledge Graph and Multi-Task Learning. *Foods* 12(10) (2023)2079.
33. D. Mckensy-Sambola, M. Á. Rodríguez-García, F. García-Sánchez, R. Valencia-García, Ontology-based nutritional recommender system. *Applied Sciences* 12(1) (2021)143.

34. C. Iwendi, S. Khan, J. H. Anajemba, A. K. Bashir, & F. Noor, Realizing an efficient IoMT-assisted patient diet recommendation system through machine learning model. *IEEE access* 8 (2020)28462-28474.
35. M. Rostami, M. Oussalah, V. Farrahi, A novel time-aware food recommender-system based on deep learning and graph clustering. *IEEE Access* 10 (2022)52508-52524.
36. R. Shandilya, S. Sharma, J. Wong, Mature-Health: Health Recommender System for Mandatory Feature choices. *arXiv preprint arXiv:2304.09099*(2023).
37. M. Rostami, U. Muhammad, S. Forouzandeh, K. Berahmand, V. Farrahi, M. Oussalah, An effective explainable food recommendation using deep image clustering and community detection. *Intelligent Systems with Applications* 16 (2022)200157.
38. Toledo, R. Y., Alzahrani, A. A., & Martinez, L. (2019). A food recommender system considering nutritional information and user preferences. *IEEE Access*, 7, 96695-96711.
39. Iwendi, C., Khan, S., Anajemba, J. H., Bashir, A. K., & Noor, F. (2020). Realizing an efficient IoMT-assisted patient diet recommendation system through machine learning model. *IEEE Access*, 8, 28462-28474.
40. Salloum, G., & Tekli, J. (2021). Automated and personalized meal plan generation and relevance scoring using a multi-factor adaptation of the transportation problem. *Soft Computing*, 1-25.
41. Chen, Y. S., Cheng, C. H., & Hung, W. L. (2021). A systematic review to identify the effects of tea by integrating an intelligence-based hybrid text mining and topic model. *Soft Computing*, 25(4), 3291-3315.
42. Masetic, Z., & Subasi, A. (2016). Congestive heart failure detection using random forest classifier. *Computer methods and programs in biomedicine*, 130, 54-64.
43. Cervantes, J., Garcia-Lamont, F., Rodríguez-Mazahua, L., & Lopez, A. (2020). A comprehensive survey on support vector machine classification: Applications, challenges and trends. *Neurocomputing*, 408, 189-215.
44. An, T. K., & Kim, M. H. (2010, October). A new diverse AdaBoost classifier. In *2010 International conference on artificial intelligence and computational intelligence* (Vol. 1, pp. 359-363). *IEEE*.
45. Bahad, P., & Saxena, P. (2020). Study of adaboost and gradient boosting algorithms for predictive analytics. In *International Conference on Intelligent Computing and Smart Communication 2019* (pp. 235-244). Springer, Singapore.
46. Bühlmann, P. (2012). Bagging, boosting and ensemble methods. In *Handbook of computational statistics* (pp. 985-1022). Springer, Berlin, Heidelberg.
47. Khan, Q. W., Iqbal, K., Ahmad, R., Rizwan, A., Khan, A. N., & Kim, D. (2024). An intelligent diabetes classification and perception framework based on ensemble and deep learning method. *PeerJ Computer Science*, 10, e1914.
48. Hooda, D., & Rani, R. (2021). An Ontology driven model for detection and classification of cardiac arrhythmias using ECG data. *Journal of Intelligent Information Systems*, 1-27.
49. Sacenti, J. A., Fileto, R., & Willrich, R. (2021). Knowledge graph summarization impacts on movie recommendations. *Journal of Intelligent Information Systems*, 1-24.
50. Sun, L., Liu, X., Liu, Y., Wang, T., Guo, L., Zheng, X., & Luo, Y. (2021). A novel deep recommend model based on rating matrix and item attributes. *Journal of Intelligent Information Systems*, 57(2), 295-319.
51. Jain, A., Singhal, A. (2022). Personalized Food Recommendation—State of Art and Review. In: Hu, YC., Tiwari, S., Trivedi, M.C., Mishra, K.K. (eds) *Ambient Communications and Computer Systems. Lecture Notes in Networks and Systems*, vol 356. Springer, Singapore. [https://doi.org/10.1007/978-981-16-7952-0\\_15](https://doi.org/10.1007/978-981-16-7952-0_15).
52. Jain, A., Singhal, A. (2023). Early Diabetes Prediction Using Deep Ensemble Model and Diet Planning. In: Devedzic, V., Agarwal, B., Gupta, M.K. (eds) *Proceedings of the International Conference on Intelligent Computing, Communication, and Information Security. ICCCIS 2022. Algorithms for Intelligent Systems*. Springer, Singapore.