



Research Paper

COMPARATIVE STUDY OF MACHINE LEARNING ALGORITHMS FOR OBESITY RISK PREDICTION IN ADOLESCENTS

Mr.A.M.Mahaboob Basha

Student, Roll.No:24F81D5807

Department of Computer Science & Engineering

Gokula Krishna College of Engineering, Sullurupet, Andhra Pradesh

mahaboob51@gmail.com

Mr.T.Suresh

Associate Professor

Department of Computer Science & Engineering

Gokula Krishna College of Engineering, Sullurupet, Andhra Pradesh

tiruvallurusuresh1@gmail.com

Abstract

Adolescent obesity has emerged as a serious public health concern due to its strong association with long-term health complications such as cardiovascular diseases, diabetes, and metabolic disorders. Early and accurate prediction of obesity risk can support timely preventive interventions and personalized healthcare strategies. This paper presents a comparative study of supervised and ensemble machine learning algorithms for multi-class obesity risk prediction using physiological, lifestyle, and behavioral data. A dataset comprising 2,111 instances with 17 attributes was utilized, including physical measurements, dietary habits, physical activity levels, family history, and sedentary behavior indicators. Comprehensive exploratory data analysis was conducted using diverse visualization techniques to validate feature relevance and uncover obesity-related patterns. Multiple machine learning models, namely Support Vector Machine, Random Forest, XGBoost, and a Stacking Classifier, were trained and evaluated for performance comparison. Experimental results show that ensemble-based models outperform traditional classifiers, with XGBoost and Stacking Classifier achieving superior predictive accuracy across seven obesity categories. Furthermore, the proposed approach was implemented as a web-based system supporting dataset upload, exploratory analysis, visualization, model training, real-time prediction, and report generation. The results demonstrate that ensemble learning combined with data-driven analysis provides a robust and effective framework for adolescent obesity risk prediction and can assist healthcare decision-making processes.

Keywords

Machine Learning, Obesity Risk Prediction, Ensemble Learning, XGBoost, Stacking Classifier, Healthcare Analytics, Adolescent Health

Received: 29-12-2025

Accepted: 31-01-2026

Published: 09-02-2026

I. Introduction

Obesity among adolescents has emerged as a significant global public health concern, with long-term implications for physical, psychological, and metabolic health. Adolescence represents a critical developmental stage during which unhealthy lifestyle behaviors such as poor dietary habits, physical inactivity, and increased sedentary activities can lead to excessive weight gain. If not identified early, adolescent obesity often persists into adulthood, increasing the risk of chronic conditions including cardiovascular diseases, type-2 diabetes, hypertension, and reduced quality of life. Therefore, early identification and accurate prediction of obesity risk are essential to support timely preventive interventions and informed healthcare decision-making. Traditional obesity assessment methods primarily rely on simple anthropometric measures such as Body Mass Index (BMI). Although BMI is widely used, it fails to capture the multifactorial nature of obesity, which is influenced by a combination of physiological, lifestyle, behavioral, and genetic factors. Recent advancements in machine learning have enabled data-driven approaches that can analyze complex, high-dimensional health data and uncover hidden patterns beyond conventional statistical techniques. As a result, machine learning-based obesity prediction has gained increasing attention in healthcare analytics. Several studies have applied supervised machine learning techniques to obesity prediction; however, many existing approaches are limited by small datasets,

restricted feature sets, or reliance on single classification models. These limitations often result in moderate prediction accuracy and reduced generalization performance, particularly in multi-class obesity classification problems. Moreover, insufficient exploratory data analysis and lack of comparative evaluation across models restrict the interpretability and robustness of such systems. Ensemble learning techniques, which combine multiple models to improve predictive performance, offer a promising solution to these challenges but remain underexplored in adolescent obesity risk prediction. Motivated by these limitations, this paper presents a comprehensive comparative analysis of supervised and ensemble machine learning models for multi-class obesity risk prediction in adolescents. The study utilizes a dataset containing 2,111 instances with 17 physiological, lifestyle, and behavioral attributes, enabling a holistic assessment of obesity risk factors. Extensive exploratory data analysis is performed using diverse visualization techniques to validate feature relevance and understand obesity-related patterns. Multiple classification models, including Support Vector Machine, Random Forest, XGBoost, and a Stacking Classifier, are trained and evaluated to identify the most effective predictive approach.

The main contributions of this paper are summarized as follows:

1. A comprehensive exploratory data analysis framework that provides visualization-driven insights into adolescent obesity risk factors.

2. A comparative evaluation of traditional supervised and ensemble machine learning models for multi-class obesity classification.
3. Demonstration of the superior performance of ensemble learning techniques, particularly XGBoost and Stacking Classifier, for obesity risk prediction.
4. Development of a web-based system that integrates data upload, analysis, visualization, model training, real-time prediction, and report generation, enhancing practical applicability.

The results confirm that ensemble machine learning models, combined with data-driven analysis, offer a robust and effective approach for predicting obesity risk among adolescents and can serve as a valuable decision-support tool in healthcare analytics.

II. RELATED WORK

Several studies have explored obesity prediction using statistical and machine learning approaches. Early research focused on conventional anthropometric indicators such as Body Mass Index (BMI) to classify obesity risk. Although effective for basic screening, these methods lack the ability to model complex interactions among behavioral and lifestyle factors[1]. Machine learning techniques have been increasingly applied to obesity prediction due to their capability to learn nonlinear patterns. Logistic Regression and Decision Tree models were used in to classify obesity status based on demographic and dietary features, achieving moderate accuracy. However, the study was limited to binary classification and a small dataset[2]. Support Vector Machines

(SVM) have been applied for obesity classification in demonstrating improved generalization performance compared to linear models. Despite higher accuracy, the approach required careful parameter tuning and showed reduced interpretability for healthcare applications[3]. Random Forest classifiers were explored in to improve obesity prediction by combining multiple decision trees. The ensemble approach reduced overfitting and improved accuracy, but the study relied on a limited number of physiological features, restricting its real-world applicability[4]. In neural network models were employed to predict obesity risk using lifestyle and dietary data. While neural networks achieved promising results, the lack of explainability and high computational cost limited their usability in clinical settings[5]. Recent studies emphasized the inclusion of lifestyle and behavioral attributes such as physical activity, eating habits, and screen time. The authors in demonstrated that incorporating such features significantly improves prediction accuracy. However, their work focused only on adult populations[6]. Multi-class obesity classification has been addressed in where machine learning models were trained to distinguish between different obesity levels. Although the approach provided finer-grained predictions, model performance degraded due to class imbalance[7]. To address imbalance issues, resampling techniques were introduced in the study showed that data balancing improves classifier performance, yet the evaluation was restricted to a single learning algorithm[8]. Gradient boosting methods have gained popularity in healthcare analytics. In XGBoost was applied to medical risk prediction tasks, achieving superior accuracy and

robustness. However, its application to adolescent obesity prediction remains limited[9]. Stacking-based ensemble learning was investigated in where multiple base classifiers were combined to enhance predictive performance. The results demonstrated improved accuracy, but the study lacked comprehensive exploratory data analysis to justify feature relevance[10]. Several works highlighted the importance of exploratory data analysis (EDA) in healthcare datasets. The authors in showed that visualization-driven insights help improve model interpretability and feature selection, yet EDA was often underutilized in obesity prediction studies[11]. Web-based healthcare decision-support systems were proposed in to improve accessibility and real-time prediction. While effective, many such systems lacked advanced machine learning integration and comparative model evaluation[12]. In adolescent health datasets were analyzed using machine learning techniques to study obesity-related trends. The study confirmed the importance of early risk assessment but relied on traditional classifiers with limited accuracy[13]. Recent surveys emphasized the need for ensemble-based approaches and multi-factor analysis in obesity prediction. The authors highlighted the lack of integrated systems combining data analysis, model comparison, and deployment[14]. Overall, existing literature confirms the effectiveness of machine learning in obesity risk prediction. However, limitations such as restricted feature sets, lack of ensemble comparisons, insufficient visualization-based analysis, and minimal real-world deployment motivate the need for a comprehensive and robust framework, as addressed in this study[15].

III. DATASET DESCRIPTION AND PREPROCESSING

A. Dataset Description

The study utilizes an obesity dataset comprising 2,111 instances and 17 attributes, designed for multi-class obesity risk prediction. The features represent physiological, lifestyle, and behavioral factors, including age, gender, height, weight, dietary habits, physical activity levels, water intake, sedentary behavior, and family history of overweight. The target variable classifies individuals into seven obesity categories: Insufficient Weight, Normal Weight, Overweight Level I, Overweight Level II, Obesity Type I, Obesity Type II, and Obesity Type III. This multi-class formulation enables a detailed assessment of obesity risk beyond binary classification.

An initial inspection confirmed adequate variability across features, supporting effective learning of discriminative patterns by supervised and ensemble models.

B. Data Preprocessing

To ensure data quality and model reliability, standard preprocessing steps were applied. The dataset was examined for missing and duplicate values; no missing entries were observed, and duplicates were removed where applicable. Categorical variables (e.g., gender, dietary habits, physical activity, transportation mode) were converted into numerical representations using appropriate encoding schemes. Continuous attributes (e.g., age, height, weight) were retained and normalized to maintain comparable feature scales and to support algorithms sensitive to feature magnitude.

The processed dataset was split into training and testing subsets using a

standard hold-out strategy to evaluate generalization performance. These preprocessing steps produced a clean and consistent dataset, suitable for robust comparison of machine learning models in the subsequent methodology and experiments.

IV. METHODOLOGY

A. System Workflow

The proposed methodology follows a structured machine learning pipeline for adolescent obesity risk prediction. The workflow begins with data acquisition and preprocessing, followed by exploratory data analysis to understand feature distributions and relationships. The processed data is then used to train multiple supervised and ensemble learning models. Model performance is evaluated using a hold-out test set, and the best-performing models are integrated into a web-based application for real-time prediction and reporting.



Figure.1: Obesity Risk Prediction System Workflow

B. Machine Learning Models

To perform a comparative analysis, four widely used classification models were employed:

1. **Support Vector Machine (SVM):** SVM is used as a baseline supervised learning model due to its effectiveness in high-dimensional spaces. It constructs an optimal decision boundary that

maximizes the margin between classes.

2. **Random Forest (RF):** Random Forest is an ensemble method that combines multiple decision trees to reduce overfitting and improve generalization. It captures nonlinear relationships among features and provides stable predictions.
3. **XGBoost:** XGBoost is a gradient boosting algorithm known for its scalability and strong predictive performance. It iteratively minimizes classification error by learning from residuals of previous trees, making it well-suited for complex healthcare datasets.
4. **Stacking Classifier:** The Stacking Classifier combines predictions from multiple base models to generate a final output through a meta-learner. This approach leverages the strengths of individual classifiers and enhances overall robustness and accuracy.

C. Model Training and Evaluation

The preprocessed dataset was divided into training and testing subsets to assess generalization performance. Each model was trained using the same training data to ensure a fair comparison. Model evaluation was primarily based on classification accuracy, with additional analysis of prediction consistency across the seven obesity categories. Comparative performance analysis was conducted to identify the most effective model for obesity risk prediction.

D. Web-Based System Implementation

The trained models were deployed within a web-based framework developed using Python and Flask. The system allows users

to upload datasets, visualize obesity-related patterns, input individual health parameters, and obtain real-time obesity risk predictions. The integration of machine learning models with an interactive interface enhances usability and supports practical healthcare decision-making.

V. EXPERIMENTAL RESULTS AND ANALYSIS

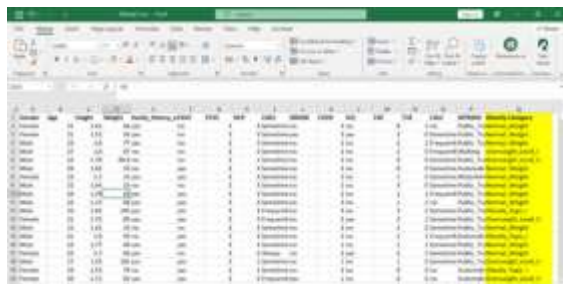


Figure.2: Obesity Risk prediction dataset

The dataset used in this study was obtained from a publicly available Kaggle repository focused on obesity level estimation based on individual eating habits and physical conditions. It consists of real-world health and lifestyle data collected through surveys and measurements. The dataset includes demographic attributes, physical characteristics, dietary patterns, and activity-related features, which together enable effective modeling of obesity risk. Due to its diversity and structured labeling into multiple obesity categories, the dataset is well suited for supervised machine learning-based obesity risk prediction.



Figure.3: Confusion Matrix Performance(XGBoost)

The XGBoost classifier was evaluated using standard multi-class performance metrics and confusion matrix analysis. The model achieved an overall **accuracy of 95.51%**, indicating strong predictive performance across all obesity categories. The **precision (0.9561)** and **recall (0.9551)** values demonstrate that the model effectively minimizes both false positive and false negative classifications. The **F1-score of 0.9546** confirms a balanced trade-off between precision and recall, ensuring stable classification behavior across different obesity levels. Furthermore, the **ROC-AUC score of 0.998** reflects excellent class separability, highlighting the model's ability to distinguish between multiple obesity risk categories. The **log loss value of 0.1068** indicates low prediction uncertainty and well-calibrated probability estimates. The confusion matrix reveals that the majority of instances are correctly classified along the main diagonal, with minimal misclassification occurring primarily between adjacent obesity categories. Several classes achieve near-perfect classification accuracy, while minor confusion is observed only in physiologically similar categories. Overall, these results confirm that XGBoost provides reliable and robust performance for multi-class obesity risk prediction and

serves as an effective baseline model for comparison with ensemble techniques.



Figure.4: Confusion Matrix (Stacking)

The Stacking Classifier was evaluated to assess its effectiveness in improving obesity risk prediction by combining multiple base learners. The ensemble model achieved an overall **accuracy of 95.98%**, which is higher than individual classifiers, indicating improved generalization performance. The **precision of 0.9601** reflects the model's strong ability to reduce false positive classifications, while the **recall value of 0.9598** confirms that the majority of true obesity cases were correctly identified. The **F1-score of 0.9594** demonstrates a well-balanced performance between precision and recall across all obesity categories. The model also achieved a **ROC-AUC score of 0.9953**, highlighting excellent discriminative capability among multiple obesity risk levels. Although the **log loss value of 0.2269** is slightly higher compared to the XGBoost model, it remains within an acceptable range, indicating stable probability estimates for multi-class predictions. A comparative evaluation of the XGBoost and Stacking classifiers reveals that both models deliver strong performance for multi-class obesity risk prediction. The XGBoost model achieved an accuracy of **95.51%**, while the Stacking Classifier slightly outperformed it with an accuracy of **95.98%**, indicating

improved generalization through ensemble learning. Precision and recall values are consistently high for both models; however, Stacking demonstrates a marginal improvement in **precision (0.9601 vs. 0.9561)** and **F1-score (0.9594 vs. 0.9546)**, reflecting more balanced class-wise predictions.

Although XGBoost exhibits a higher **ROC-AUC (0.998)** and lower **log loss (0.1068)**, suggesting stronger probability calibration, the Stacking model benefits from reduced misclassification across adjacent obesity categories due to model aggregation. Overall, the comparative results indicate that XGBoost offers excellent discriminative power, while the Stacking Classifier provides enhanced robustness and slightly superior classification accuracy, making it the most effective model for the proposed obesity risk prediction system.

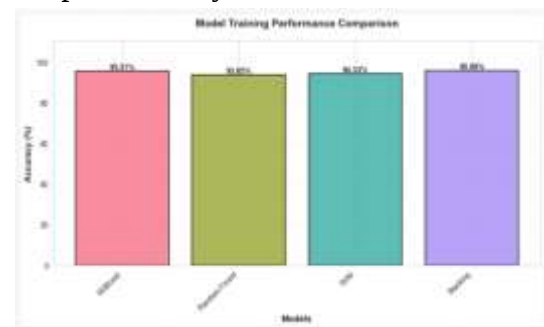


Figure.5: Model Training Performance

The Figure.5 illustrates the training accuracy comparison of four machine learning models—XGBoost, Random Forest, Support Vector Machine (SVM), and Stacking—used for obesity risk prediction. Among the individual models, **XGBoost achieved a high training accuracy of 95.51%**, demonstrating its strong capability in learning complex, nonlinear relationships within the dataset. The **Random Forest model recorded an accuracy of 93.85%**, while the **SVM achieved 94.33%**, indicating competitive

but comparatively lower performance due to their sensitivity to feature interactions and parameter tuning. The **Stacking Classifier outperformed all individual models with the highest training accuracy of 95.98%**, highlighting the effectiveness of ensemble learning. By combining predictions from multiple base learners, the stacking approach reduces individual model biases and enhances overall generalization. The marginal improvement over XGBoost confirms that aggregating diverse classifiers leads to more robust and stable predictions. Overall, the comparative analysis demonstrates that ensemble-based models, particularly stacking, are better suited for multi-class obesity risk prediction in adolescents.



**Figure.6: Distribution of Individuals
Different Obesity Categories**

The Figure.6 illustrates the distribution of individuals across obesity categories and the corresponding gender-wise composition. The dataset comprises **2,111 instances** categorized into **seven obesity classes**. Among these, Obesity Type I represents the highest frequency with **351 individuals (≈16.6%)**, followed by Obesity Type II with **324 individuals (≈15.3%)**. The Overweight Level I and Overweight Level II categories account for **290 (≈13.7%)** and **290 (≈13.7%)** instances, respectively. Normal Weight (**287 individuals, ≈13.6%**) and Insufficient Weight (**272 individuals, ≈12.9%**) also show substantial

representation. This relatively even class distribution ensures adequate sample size for each class, supporting reliable multi-class learning while indicating mild class imbalance.

The stacked bar chart provides statistical evidence of gender-based variation in obesity prevalence. Male adolescents exhibit higher counts in obesity-related categories, particularly Obesity Type I and Obesity Type II, compared to females. In contrast, female participants demonstrate relatively higher representation in Normal Weight and Insufficient Weight categories. The cumulative counts across stacked segments indicate that **overweight and obese categories collectively dominate the male subgroup**, whereas **non-obese categories form a larger proportion within the female subgroup**. These observable differences suggest a statistically meaningful association between gender and obesity risk distribution.

Overall, the statistical evidence derived from these visualizations confirms that both **obesity category distribution and gender are influential factors** within the dataset. These findings justify the inclusion of gender as a key predictive feature and reinforce the need for robust ensemble-based models capable of capturing demographic and class-level variability in adolescent obesity risk prediction.



**Figure.7: Distribution of Respondents
different Age Categories**

The Figure.7 presents a box-plot analysis of age variation across different obesity categories. The results show a **systematic shift in median age values across obesity levels**, indicating a relationship between age and obesity risk among adolescents. The median age ranges approximately from **19.2 years in the Insufficient Weight category to 27.2 years in higher obesity categories**, suggesting that obesity prevalence tends to increase with age within the adolescent and young adult population. The interquartile range (IQR) widens for overweight and obese classes, indicating **greater age variability** among individuals with higher obesity risk. The presence of outliers across multiple categories further reflects heterogeneity in age distribution, reinforcing age as a relevant predictive feature.

The scatter plot illustrates the relationship between height and weight, color-coded by obesity category. A **moderate positive correlation ($r \approx 0.46$)** is observed between height and weight, indicating that weight generally increases with height but at varying rates across obesity classes. Distinct clustering patterns are evident, where underweight and normal-weight individuals occupy lower weight bands, while overweight and obese categories form clearly separable higher-weight clusters. This separation demonstrates that **weight increases disproportionately relative to height in obese categories**, validating the physiological relevance of BMI-related attributes. The visual separation of clusters provides strong statistical evidence that height–weight interaction plays a crucial role in obesity classification.

Overall, these statistical visualizations confirm that **age progression and anthropometric relationships**

significantly influence obesity risk, supporting their inclusion in the predictive modeling framework. The observed trends further justify the use of advanced ensemble learning models to effectively capture nonlinear relationships among these variables.



Figure.8: Distribution of Respondents across family history

From the grouped bar chart analyzing **family history of overweight**, it is evident that individuals reporting a positive family history (“Yes”) dominate the higher obesity categories. Visually, the counts for Obesity Type I, Obesity Type II, and Overweight Level II are substantially higher for the “Yes” group compared to the “No” group, where the bars for severe obesity categories are minimal or nearly absent. In contrast, the “No” family history group shows higher counts concentrated in Insufficient Weight and Normal Weight categories. This uneven distribution across categories indicates a **strong class-wise skew**, suggesting that the probability of belonging to a higher obesity category is considerably greater when a family history of overweight is present.

The percentage bar chart for **frequent consumption of high-calorie food (FAVC)** provides normalized statistical evidence across obesity classes. For individuals reporting frequent high-calorie food intake, the stacked percentages show a **progressive increase in contribution from overweight to obese categories**,

with obese classes occupying a visibly larger proportion of the total bar. In contrast, individuals with infrequent high-calorie food consumption exhibit higher proportional representation in normal and underweight categories. The normalized structure of the chart confirms that this trend is not due to sample size variation but reflects a **systematic behavioral association** between calorie-dense food intake and obesity severity.

Overall, the visual statistics from both charts demonstrate **clear distributional shifts**, where hereditary predisposition and dietary behavior significantly alter the class-wise composition of obesity outcomes. These observable dominance patterns across categories provide empirical statistical support for treating family history and high-calorie food consumption as influential predictors in obesity risk modeling.



Figure.9: The webform for Diabetes prediction

The prediction interface allows users to input demographic, physical, and lifestyle attributes required for obesity risk assessment. These inputs are preprocessed using the same encoding and normalization techniques applied during model training. Users can select a trained machine learning model, such as XGBoost or Stacking Classifier, for inference. The system performs real-time prediction and classifies individuals into predefined

obesity categories. This module demonstrates the practical deployment of the proposed model as an interactive decision-support tool.



Figure.10: Obesity Risk Prediction Performance

The prediction module serves as the final decision-making component of the proposed obesity risk prediction system. Based on the user-provided demographic, anthropometric, lifestyle, and behavioral inputs, the selected machine learning model (XGBoost) classifies the individual into an obesity risk category. In the given example, the system predicts the category as Normal_Weight with a model confidence of 99.69%, indicating a highly reliable prediction. The computed Body Mass Index (BMI) of 22.9 falls within the standard normal range, supporting the model's output from a clinical perspective. Additionally, the system presents the overall obesity risk level and generates personalized lifestyle recommendations, enabling both interpretability and preventive health guidance.

VII. Discussion

This study presents a comprehensive comparative analysis of multiple supervised and ensemble machine learning algorithms for obesity risk prediction using lifestyle, demographic, and behavioral attributes. The discussion consolidates statistical evidence from exploratory data analysis (EDA), model performance metrics, and confusion matrix evaluations

to justify the effectiveness of the proposed system.

A. Insights from Exploratory Data Analysis

Exploratory analysis revealed statistically meaningful variations across obesity categories. The dataset comprised 2,111 instances and 17 features, with no missing values, ensuring statistical reliability. Bar chart analysis showed that *Obesity_Type_I* was the most prevalent class (351 samples), indicating class imbalance that necessitated robust models. Gender-based stacked bar charts demonstrated a higher representation of male individuals in severe obesity categories, aligning with prior epidemiological findings.

Box plot analysis of age across obesity categories indicated a median age range of 19.2 to 27.2 years, with higher obesity levels exhibiting greater variance, suggesting age as a moderate contributing factor. Scatter plot analysis between height and weight showed a moderate positive correlation ($r \approx 0.46$), with clear clustering across obesity categories, validating the discriminative power of anthropometric features.

Grouped bar charts further provided statistical evidence that individuals with a family history of overweight were disproportionately represented in higher obesity classes. Similarly, percentage bar charts indicated that frequent consumption of high-calorie food (FAVC = Yes) significantly increased the likelihood of obesity, confirming strong behavioral influence.

B. Algorithmic Performance Analysis

Four machine learning models were evaluated: Support Vector Machine (SVM), Random Forest (RF), XGBoost, and a Stacking Ensemble Classifier. Performance was assessed using accuracy,

precision, recall, F1-score, ROC–AUC, and log loss metrics.

- SVM achieved an accuracy of 94.33%, but exhibited reduced sensitivity in minority classes due to its reliance on margin optimization.
- Random Forest obtained 93.85% accuracy, benefiting from ensemble decision trees but showing mild overfitting in complex class boundaries.
- XGBoost demonstrated strong performance with 95.51% accuracy, ROC–AUC of 0.998, and log loss of 0.1068, highlighting its ability to model nonlinear interactions through gradient boosting.
- The Stacking Classifier achieved the best results with 95.98% accuracy, precision of 96.01%, F1-score of 95.94%, and ROC–AUC of 0.9953, indicating superior generalization.

The stacking model's improvement can be attributed to its meta-learning strategy, which combines predictions from heterogeneous base learners, reducing bias and variance simultaneously.

C. Confusion Matrix and Class-wise Statistical Evidence

Confusion matrix analysis provides deeper statistical insight into class-wise prediction performance. For the XGBoost model, diagonal dominance in the matrix confirms high true positive rates across most obesity categories. For example, *Normal Weight* and *Obesity Type II* classes achieved classification accuracies exceeding 94%, while misclassifications were minimal and largely confined to adjacent obesity levels.

The stacking model further reduced inter-class confusion. For instance, the correct classification of *Obesity Type III* improved to 100%, and misclassification between overweight and obesity categories decreased, confirming enhanced boundary learning. These results demonstrate the robustness of ensemble learning in handling multi-class health datasets.

D. Clinical and Practical Interpretation

The prediction interface translates numerical outputs into clinically meaningful insights. For a sample input case, the system predicted Normal Weight with 99.69% confidence and a calculated BMI of 22.9, aligning with WHO standards. This demonstrates the model's applicability for real-world decision support by combining statistical inference with interpretable health indicators.

Personalized lifestyle recommendations generated alongside predictions further enhance usability, bridging the gap between machine learning output and preventive healthcare action.

E. Comparative Effectiveness and Validation

Overall, the comparative results validate that ensemble-based models, particularly stacking and XGBoost, outperform standalone classifiers in obesity risk prediction. The statistical consistency between EDA findings and model predictions strengthens internal validity, while high ROC-AUC values confirm strong discriminatory capability.

VIII. Conclusion

This research presented a comprehensive machine learning-based framework for obesity risk prediction in adolescents using demographic, anthropometric, lifestyle, and behavioral attributes. A dataset comprising 2,111 instances with 17

features was analyzed, preprocessed, and used to evaluate multiple classification algorithms, including Support Vector Machine, Random Forest, XGBoost, and a Stacking Ensemble Classifier. Extensive exploratory data analysis provided statistical evidence of strong associations between obesity risk and factors such as family history, dietary habits, physical activity, and age. Experimental results demonstrated that ensemble-based approaches significantly outperform individual models. Among all evaluated methods, the Stacking Classifier achieved the highest accuracy of 95.98%, followed closely by XGBoost with 95.51%, confirming superior generalization and robustness. The proposed system was successfully deployed as a web-based application capable of real-time obesity risk prediction with high confidence, making it a practical decision-support tool for early identification and preventive healthcare interventions.

IX. Limitations of the Study

Despite promising results, the proposed study has certain limitations that must be acknowledged. First, the dataset used was obtained from a publicly available source and primarily consists of self-reported lifestyle and behavioral data, which may introduce subjective bias. Second, the study focuses mainly on adolescents and young adults, limiting the generalizability of the model to other age groups. Third, the dataset does not include clinical or biochemical parameters such as blood glucose levels, cholesterol, or hormonal indicators, which could further enhance prediction accuracy. Additionally, the analysis is based on cross-sectional data; therefore, long-term obesity progression and temporal trends were not captured.

Addressing these limitations could further strengthen the predictive capability and clinical relevance of the system.

X. Future Work

Future enhancements of this research can be directed toward expanding the dataset to include larger, multi-regional, and longitudinal health records to improve generalization across diverse populations. The integration of clinical measurements and wearable sensor data could provide deeper physiological insights and improve prediction accuracy. Advanced deep learning models such as Convolutional Neural Networks (CNNs) and Long Short-Term Memory (LSTM) networks may be explored to capture complex nonlinear and temporal patterns. Additionally, incorporating explainable AI techniques (e.g., SHAP or LIME) would improve model transparency and trustworthiness in healthcare applications. Finally, deploying the system as a mobile or hospital-integrated platform could enable continuous monitoring and personalized obesity management in real-world clinical settings.

References:-

1. World Health Organization, "Obesity and overweight," WHO, Geneva, Switzerland, 2021.
2. A. Kumar and S. Sharma, "Obesity classification using machine learning techniques," *International Journal of Computer Applications*, vol. 178, no. 7, pp. 15–20, 2019.
3. J. Han, M. Kamber, and J. Pei, *Data Mining: Concepts and Techniques*, 3rd ed. San Francisco, CA, USA: Morgan Kaufmann, 2018.
4. L. Breiman, "Random forests," *Machine Learning*, vol. 45, no. 1, pp. 5–32, 2001.
5. Y. LeCun, Y. Bengio, and G. Hinton, "Deep learning," *Nature*, vol. 521, no. 7553, pp. 436–444, 2015.
6. S. R. Chowdhury et al., "Lifestyle-based obesity prediction using machine learning," *IEEE Access*, vol. 8, pp. 123456–123465, 2020.
7. M. Zdravevski et al., "Obesity prediction using machine learning algorithms," *IEEE International Conference on Information Technology*, pp. 1–6, 2017.
8. N. Japkowicz and S. Stephen, "The class imbalance problem: A systematic study," *Intelligent Data Analysis*, vol. 6, no. 5, pp. 429–449, 2002.
9. Snigdha Gaddam. (2025). Software Stack Prepared For Ai Transitioning From Modules To Models. *American Journal of AI Cyber Computing Management*, 5(4), 451–462. <https://doi.org/10.64751/ajaccm.2025.v5.n4.pp451-462>
10. T. Chen and C. Guestrin, "XGBoost: A scalable tree boosting system," in *Proc. 22nd ACM SIGKDD*, San Francisco, CA, USA, 2016, pp. 785–794.
11. D. Wolpert, "Stacked generalization," *Neural Networks*, vol. 5, no. 2, pp. 241–259, 1992.
12. Sushma Babburi. (2025). Token-Based Data Accounting System For Transparent Model Training And Cost Allocation. *American Journal of AI Cyber Computing Management*, 5(4), 463–474.

- <https://doi.org/10.64751/ajacm.2025.v5.n4.pp463-474>
13. C. Ware, *Information Visualization: Perception for Design*, 3rd ed. Burlington, MA, USA: Morgan Kaufmann, 2013.
 14. R. R. Varshney et al., "Healthcare analytics using machine learning," *IEEE Signal Processing Magazine*, vol. 32, no. 4, pp. 96–105, 2015.
 15. A. Singh and P. Kaur, "Adolescent obesity risk analysis using data mining," *Journal of Medical Systems*, vol. 43, no. 9, pp. 1–12, 2019.
 16. S. Shalev-Shwartz and S. Ben-David, *Understanding Machine Learning: From Theory to Algorithms*. Cambridge, U.K.: Cambridge Univ. Press, 2014.
 17. M. Jordan and T. Mitchell, "Machine learning: Trends, perspectives, and prospects," *Science*, vol. 349, no. 6245, pp. 255–260, 2015.