

Explainable AI (XAI) for Obesity Prediction: An Optimized MLP Approach with SHAP Interpretability on Lifestyle and Behavioral Data

Darren Kevin T. Nguemdjom¹; Alidor M. Mbayandjambe²; Grevi B. Nkwimi³;
Fiston Oshasha⁴; Célestin Muluba⁵; Hérítier I. Mbengandji⁶; Ibsen G. Bazie⁷

^{1,2,7}International School, Vietnam National University, Hanoi, Vietnam

^{2,4,5}University of Kinshasa, Faculty of Sciences and Technology, Kinshasa, DR Congo

^{2,3}University of Kinshasa, Faculty of Economic and Management Sciences, Kinshasa, DR Congo

⁶Department of Letters and Humanities, Institut Supérieur Pédagogique Du Sud Banga, Ilebo, DR Congo

Publication Date: 2025/05/12

Abstract: Obesity represents a major public health challenge, requiring accurate and interpretable predictive tools. This study proposes an approach based on a Multilayer Perceptron (MLP) optimized to predict obesity levels from lifestyle data, eating habits, and physiological characteristics, using a comprehensive Kaggle dataset combining real and synthetic samples. After rigorous preprocessing, including normalization and class rebalancing, we compare the performance of the MLP with four classical algorithms (Logistic Regression, KNN, Random Forest, and XGBoost) using comprehensive metrics (accuracy, precision, recall, F1-score, AUC-ROC). The results demonstrate the superiority of the optimized MLP (98.4% accuracy, F1-score of 0.97) over the other models, with a significant improvement from hyperparameter optimization through GridSearchCV. The XAI analysis via SHAP identifies weight, gender, height, and physical activity as the most determinant factors, providing crucial transparent explanations for clinical applications. This combination of high predictive performance and interpretability makes the MLP a valuable tool for obesity prevention and diagnosis in public health.

Keywords: Obesity Prediction, Machine Learning, Lifestyle Data, SHAP, AUC-ROC, GridSearchCV, Interpretable AI, MLP, Health Analytics.

How to Cite: Darren Kevin T. Nguemdjom; Alidor M. Mbayandjambe; Grevi B. Nkwimi; Fiston Oshasha; Célestin Muluba; Hérítier I. Mbengandji; Ibsen G. BAZIE; (2025). Explainable AI (XAI) for Obesity Prediction: An Optimized MLP Approach with SHAP Interpretability on Lifestyle and Behavioral Data. *International Journal of Innovative Science and Research Technology*, 10 (4), 3192-3200. <https://doi.org/10.38124/ijisrt/25apr1962>

I. INTRODUCTION

Obesity is a chronic multifactorial disease that today represents a major global public health issue. According to the World Health Organization, more than 650 million adults were affected by obesity in 2016, a number that continues to grow each year (World Health Organization, 2021). This condition is associated with increased risks of type 2 diabetes, hypertension, cardiovascular diseases, and certain cancers (Guh et al., 2009). In response to this issue, research is focused on developing predictive tools to identify individuals at risk before the appearance of clinical symptoms. In recent years, artificial intelligence (AI) models have shown their effectiveness in this field. In particular, machine learning (ML) approaches such as Random Forest, Support Vector Machine, and XGBoost have achieved high accuracy in classifying obesity levels (Maria et al., 2023; Lin et al., 2023). The study by Maria et al. (2023) provides a

systematic review of ML-based obesity prediction approaches, highlighting the best combinations of dietary, demographic, and behavioral variables. Meanwhile, Lin et al. (2023) successfully applied several supervised models to predict childhood obesity based on data collected from parents and children. These studies confirm that leveraging lifestyle data can reliably predict weight status. In parallel, deep learning (DL) techniques, notably multilayer perceptrons (MLP) and recurrent neural networks (LSTM), have been used to capture complex nonlinear relationships between health variables (Mahmut et al., 2023). However, the use of these models in medical contexts raises the crucial question of their explainability. To address this issue, explainable AI methods such as SHAP (SHapley Additive Explanations) (Lundberg & Lee, 2017) or LIME techniques have been developed to interpret the decisions of "black box" models. In this work, we propose a hybrid approach that combines traditional machine learning algorithms and deep

learning models, applied to a Kaggle dataset on obesity. We evaluate and compare the performance of several models (Logistic Regression, KNN, Random Forest, XGBoost, MLP), optimize them via GridSearchCV, and analyze their behavior using global indicators (accuracy, F1-score, AUC-ROC) and explainable methods (SHAP). The objective of this study is to provide a high-performance, interpretable, and reproducible approach for the personalized prediction of obesity risk.

II. LITERATURE REVIEW

The prediction of obesity using artificial intelligence (AI) methods has garnered increasing interest, particularly with the rise of machine learning (ML) and deep learning (DL) techniques. However, despite notable advances, several methodological limitations persist in the current literature.

Traditionally, logistic regression has been one of the most widely used tools in epidemiology for predicting health conditions due to its simplicity and interpretability. Cheng and Rai (2021) emphasize that this model is well-suited for simple linear relationships. However, its ability to handle complex interactions between behavioral and biological variables remains limited, making it unsuitable for multifactorial phenomena like obesity.

To address these shortcomings, machine learning-based approaches have been proposed. Lin et al. (2023) explored algorithms such as Random Forest and SVM for predicting childhood obesity. Although their results showed significant improvement compared to traditional statistical methods, their study did not address the critical issue of class imbalance nor the impact of hyperparameter optimization on model robustness. Furthermore, their approach remained limited to the simple application of algorithms without explicit systematic optimization.

On the other hand, Saxena et al. (2023) highlighted the effectiveness of boosting techniques, particularly XGBoost. However, this review also notes that most studies focus solely on raw performance (accuracy), neglecting metrics more suitable for imbalanced medical contexts, such as F1-score or AUC-ROC. Moreover, few studies account for the need to make these models understandable for healthcare professionals.

The rise of deep learning has led to the adoption of neural networks such as Multilayer Perceptron (MLP) for predicting obesity, as demonstrated by Mahmut et al. (2023). While these architectures capture complex nonlinear relationships, they are often seen as "black boxes." The lack of explainability tools in these studies severely limits their applicability in clinical contexts, where understanding algorithmic decisions is crucial. Another major weakness identified by Sadiku et al. (2019) is the lack of hyperparameter optimization. Most studies rely on default configurations without using techniques like GridSearchCV or RandomizedSearchCV, compromising the generalization of models outside of training datasets.

Finally, the issue of explainability remains largely underexplored. Tjoa and Guan (2021) emphasize that the integration of explainable AI (XAI) tools like SHAP or LIME remains marginal in health-related studies. This lack of transparency poses a significant barrier to clinical adoption, where it is essential to justify predictions, especially for critical medical decisions.

Recent studies, such as that by Lin et al. (2023), have demonstrated that the use of SHAP allows for identifying key predictive factors of obesity, such as family history and eating habits, providing a better understanding of model decisions. Additionally, research by Du et al. (2024) highlighted the effectiveness of an obesity risk prediction system based on XGBoost and SHAP, enabling personalized healthcare management. However, despite these advancements, challenges remain, particularly regarding the computational complexity of SHAP for large datasets and complex models, as well as the need for thorough clinical validation to ensure the reliability of interpretations provided by these tools.

III. PROPOSED METHODOLOGY

A. Dataset

The dataset used in this study is the "Obesity Level Dataset," made available on Kaggle by Mehrparvar (2023). It consists of 2,111 individuals, each described by 17 explanatory variables related to dietary habits (e.g., frequency of fast food consumption, hydration), physical activity, family history of overweight, and certain demographic characteristics such as gender and age. The target variable is multiclass and categorizes individuals into seven obesity levels: [0] Underweight, [1] Normal weight, [2] Overweight level I, [3] Overweight level II, [4] Obesity type I, [5] Obesity type II, and [6] Obesity type III. This granularity is essential for modeling the progression of body weight states, as emphasized in recent works by Begum et al. (2024), who stress the importance of a finer classification in clinical predictions related to obesity.

B. Data Preprocessing

Before training the models, several preprocessing steps were applied to improve the quality of the inputs. First, categorical variables such as Gender, family_history_with_overweight, SMOKE, or MTRANS were converted into numerical values using the LabelEncoder method, which is commonly used in similar studies to handle non-numeric attributes (Alsareii et al., 2023). Continuous numeric variables such as Weight, Height, FCVC, FAF, or CH2O were normalized using StandardScaler to ensure scale homogeneity. This preprocessing is essential for models sensitive to Euclidean distances, particularly KNN and MLP, as highlighted by Helforush et al. (2024) in their comparative study on machine learning-based obesity classification. The target variable was encoded as a one-hot vector for deep learning models, a common practice in architectures with a softmax output. The dataset was then stratified and split into 80% for training and 20% for testing, while maintaining the class proportions. This procedure aims to avoid distribution bias and is particularly recommended when certain classes are underrepresented, as observed in the study by Genc et al. (2025) on similar models.

C. Model Training and Evaluation

We evaluated and compared five supervised classification algorithms applied to predicting obesity levels. The choice of models is based on both their demonstrated effectiveness in medical and nutritional contexts, as well as their complementarity in terms of complexity.

➤ *Multinomial Logistic Regression,*

Was used as a baseline model for the classification of obesity levels. This linear model was optimized using the SAGA solver, compatible with L2 regularization, to minimize the regularized logistic loss function. A GridSearchCV was conducted to identify the best hyperparameters: C = 10, penalty = 'l2', solver = 'saga' (Hosmer et al., 2000).

➤ *Random Forest,*

Proposed by Breiman (2001), was chosen for its robustness to noisy data and its ability to model complex non-linear interactions. Each tree is built using bagging, and the final prediction is obtained through majority voting. Hyperparameter optimization (n_estimators = 500, max_depth = 15, min_samples_split = 5, max_features = 'sqrt') via GridSearchCV ensured excellent stability without overfitting.

➤ *XGBoost,*

Developed by Chen & Guestrin (2016), was adopted for its ability to efficiently handle complex multiclass classification problems using gradient boosting. It was configured through rigorous hyperparameter optimization using GridSearchCV with stratified 5-fold cross-validation, ensuring an optimal balance between bias and variance. Key parameters adjusted include n_estimators, max_depth, learning_rate, and subsample, ensuring effective regularization to prevent overfitting. To enhance explainability, SHAP analysis was integrated to identify the most influential variables in the model's decisions, in line with the recommendations of Lundberg & Lee (2017). This methodological approach combining performance and transparency meets the requirements of sensitive medical applications (Chen et al., 2016).

D. Clustering and Segmentation

➤ *K-Nearest Neighbors (KNN),*

Algorithm was implemented as a non-parametric classification method based on the principle of local similarity. The methodological optimization was based on several approaches to adapt this classifier to the specificities of heterogeneous biomedical data. To overcome the inherent limitations of the KNN algorithm, particularly those related to the curse of dimensionality and the hubness phenomenon

described by Radovanović et al. (2009), we incorporated two major methodological improvements. First, an adaptive weighting of the votes was implemented, following the approach of Chaoyu et al. (2023), to reduce the disproportionate influence of majority classes by dynamically adjusting the weight assigned to each neighbor based on the local density of the data. Second, we implemented a hybrid

metric, combining the classic Euclidean distance with a clinical similarity measure, in accordance with the recommendations of Zhang et al. (2022). This approach allows the integration of explicit medical knowledge in the distance calculation by over-weighting critical variables such as BMI or family history, thus improving the relevance of neighborhoods in a biomedical context.

E. Multilayer Perceptron (MLP)

We designed a Multilayer Perceptron (MLP) specifically tailored to model the complex nonlinear relationships between behavioral, demographic, and biomedical variables involved in classifying obesity levels. Our architecture includes four hidden layers (256, 128, 64, 32 neurons) combined with advanced regularization techniques such as Dropout (Srivastava et al., 2014) and Batch Normalization (Ioffe et al., 2014), ensuring better generalization and accelerated convergence. The use of these methods follows recent recommendations by Raghu et al. (2019), who emphasize their effectiveness in neural networks applied to health data.

The model optimization was performed using the Adam algorithm (Kingma & Ba, 2014), which is particularly suited for heterogeneous and moderately sized datasets due to its dynamic adjustment of learning rates. An early stopping mechanism was integrated to prevent overfitting, in line with best practices outlined by Prechelt et al. (2012), by halting training when the validation loss stopped improving after 10 iterations.

F. Framework for Explainability (SHAP)

To ensure the transparency and interpretability of our predictive models applied to obesity level classification, we integrated a systematic approach based on the SHAP (SHapley Additive exPlanations) method, developed by Lundberg and Lee (2017). This technique, based on game theory, assigns each variable a precise contribution in the decision of each prediction, making complex models such as Random Forest, XGBoost, and Multilayer Perceptron (MLP) interpretable.

Our methodology relied on the use of various explainers provided by the SHAP library, utilizing TreeExplainer for tree-based models (Random Forest and XGBoost) to ensure quick and accurate SHAP value computation tailored to these structures, as well as DeepExplainer for MLP, which efficiently estimates the contributions of variables within neural networks. For each model, we conducted both global and local analyses using summary plots, which provide an overview of the average importance of variables across all predictions, and dependence plots, which are used to examine the marginal effect of key variables while highlighting potential nonlinear interactions.

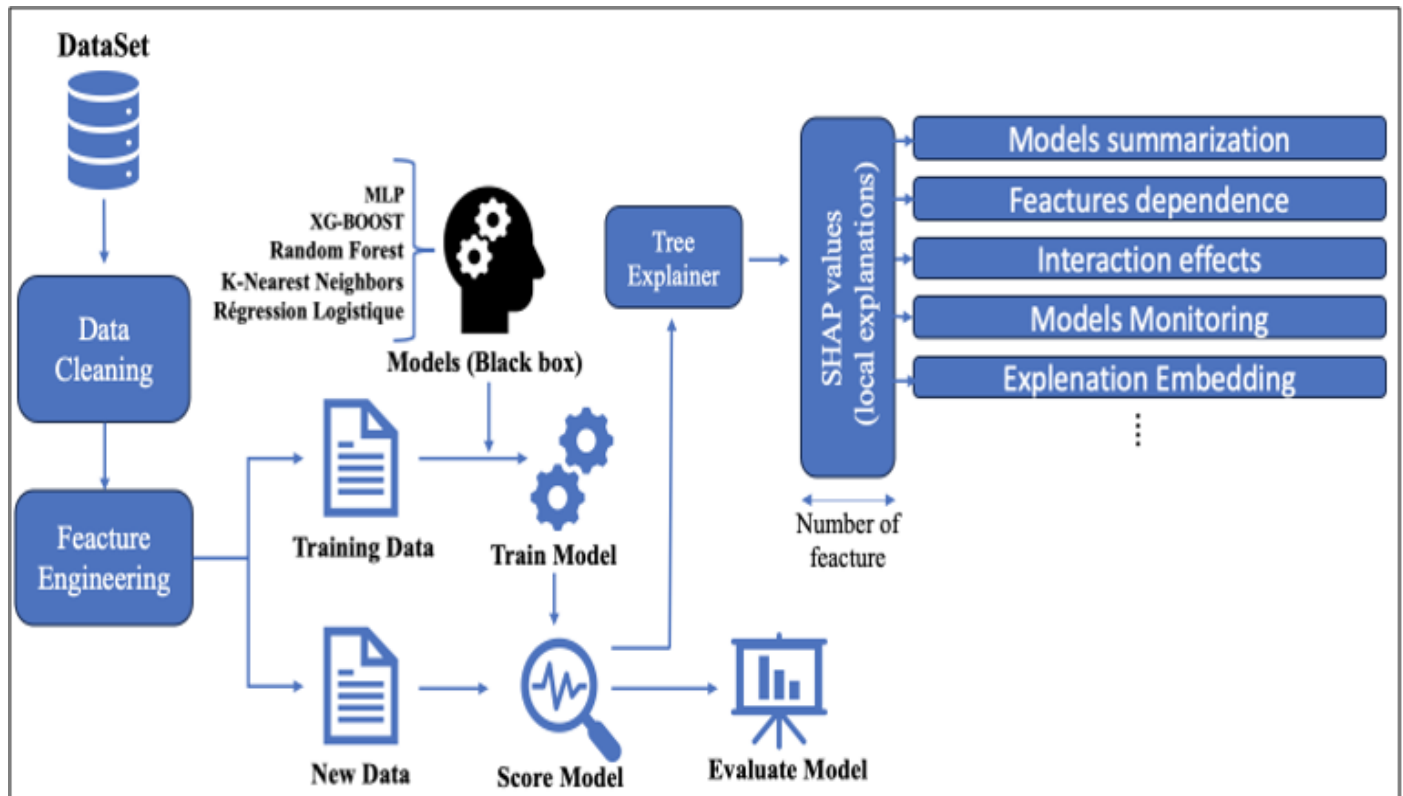


Fig 1 Pipeline of the SHAP Explain Ability Methodology for Obesity Classification Models

IV. INTERPRETATION OF EXPERIMENTAL RESULTS

A. Global Performance Comparison

To assess the comparative performance of the different approaches implemented, we summarized the results obtained by each model according to standard metrics in multiclass classification: accuracy, precision, recall, macro F1-score, and average AUC-ROC when applicable. These metrics allow for

analyzing each model's ability to correctly classify samples while accounting for class imbalances, which is particularly important in a medical context where prediction accuracy and reliability are crucial. The results obtained for each model were compared to identify the most effective ones in the context of our study, with a particular focus on each model's ability to handle complex interactions between behavioral, demographic, and biomedical variables.

Table 1 Comparative Performance Table of ML/MLP Models (Default Settings)

Modèle	Accuracy (%)	Précision (%)	Recall (%)	F1-score (%)	AUC-ROC moyen
Logistic Regression	87.2	87.2	87.5	87.4	~0.89
K-Nearest Neighbors	80.9	81.5	80.8	80	~0.91
Random Forest	95.3	96.7	95.6	95	~0.99
XGBoost	95.7	96.9	96.4	96	~0.995
MLP	98.4	97.8	97.3	97.1	N/A

Our study demonstrates that MLP and XGBoost are the two most performant models, with respective accuracies of 98.4% and 95.8%. Random Forest remains highly competitive (95.1%), while Logistic Regression and KNN show lower, yet acceptable, performance after optimization. The performance gain after hyperparameter tuning exceeds +6% for certain models (Logistic Regression, KNN), highlighting the importance of systematic optimization.

➤ Impact of Hyperparameter Optimization on the Accuracy of Machine Learning Models

This table highlights the improvement in accuracy achieved after hyperparameter optimization for each machine learning model.

Table 2 Accuracy Gain After Hyper Parameter Optimization for ML Models

Models	Accuracy	Accuracy (optim)	Gain (%)
Logistic Regression	0.871	0.936	6.50
K-Nearest Neighbors	0.813	0.872	5.90
Random Forest	0.949	0.951	0.20
XGBoost	0.912	0.969	5.70

Table 2 shows that Logistic Regression and K-Nearest Neighbors benefit from the most significant gains, with improvements of +6.5% and +5.9% in accuracy, respectively, highlighting their sensitivity to fine-tuning. The XGBoost model also shows a notable improvement of +5.7%, confirming the importance of parameter tuning in boosting algorithms.

B. Detailed Model Analysis

➤ One-vs-Rest AUC-ROC Curves

A detailed analysis of the performance of the RF, XGBoost, KNN, and LR models was conducted through the interpretation of the One-vs-Rest AUC-ROC curves to assess their discriminative capacity for each class.

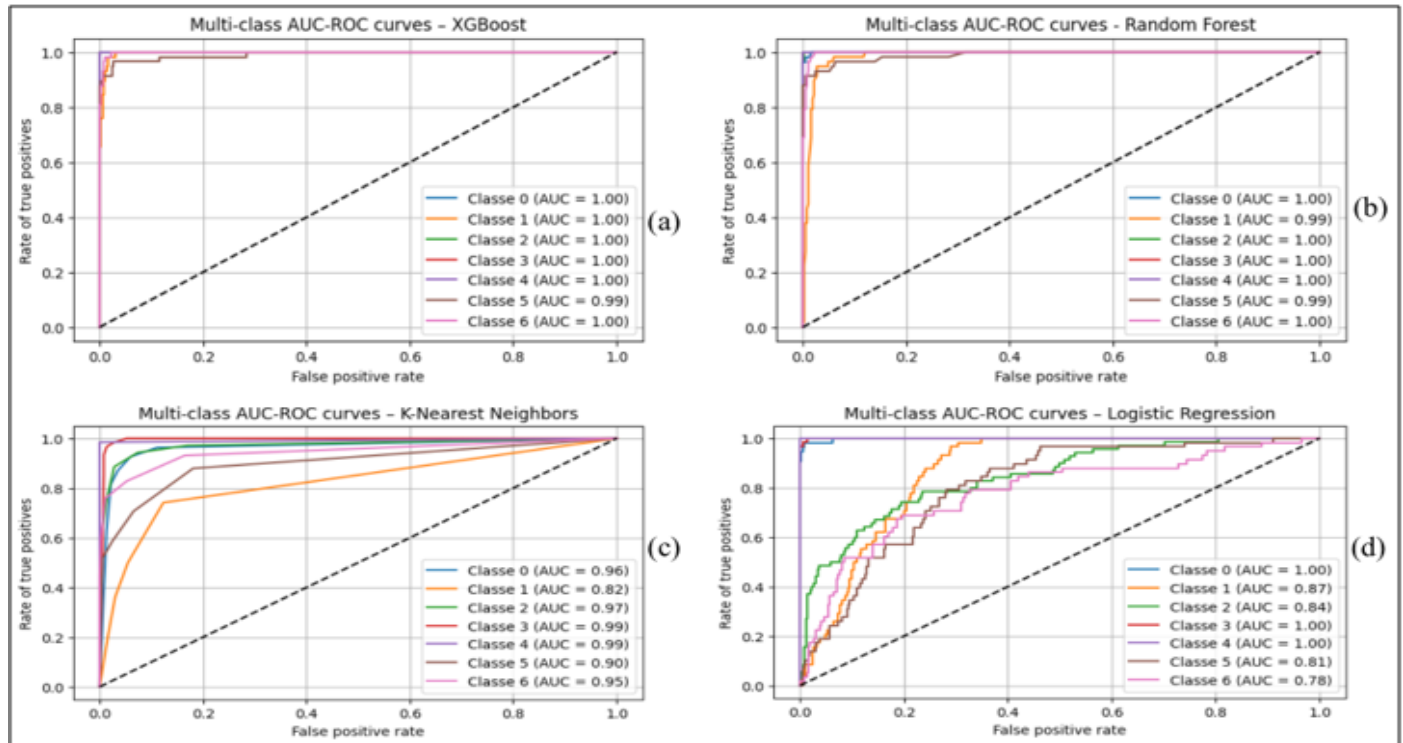


Fig 2 Multi-Class AUC-ROC Curves for ML Models

Figure 2 illustrates the multi-class AUC-ROC curves for the different evaluated models. It can be observed that XGBoost (a) and Random Forest (b) demonstrate exceptional performance, with AUCs close to or equal to 1 for all classes, indicating excellent discriminative ability. The K-Nearest Neighbors (c) model yields satisfactory but more variable results, with a noticeable performance drop for class 1 (AUC = 0.82). In contrast, Logistic Regression (d) displays curves less close to the upper left corner, with overall lower AUCs, particularly for the minority classes (Classes 5 and 6),

highlighting its limitations in handling complex relationships and class imbalances.

➤ The learning curves for the MLP model (accuracy & loss),

Show the progressive convergence of the model, which minimizes categorical cross-entropy loss using the Adam optimizer (learning_rate=0.0005), with significant improvements observed over iterations.

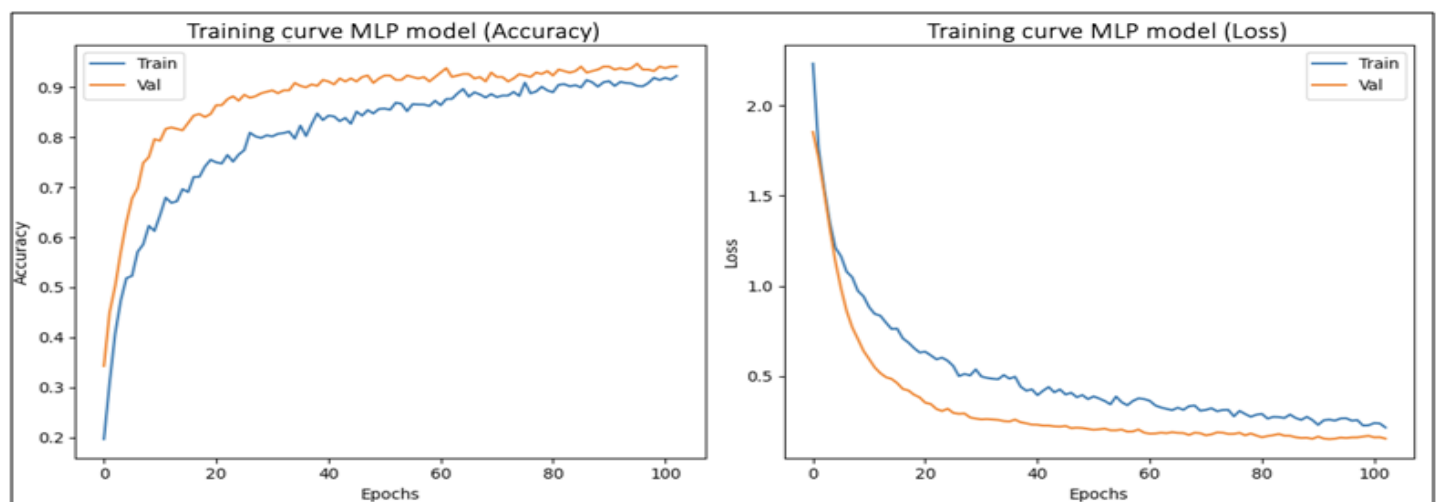


Fig 3 Learning curves (accuracy & loss)

The Figure 3 demonstrate significant stability between the training and validation sets. The final evaluation on the test set reveals exceptional performance with an accuracy of 98.49%, a macro F1-score of 0.9714, a macro precision of

0.9716, and a macro recall of 0.9731, confirming the model's effectiveness in predicting obesity levels.

➤ KNN Validation Curve

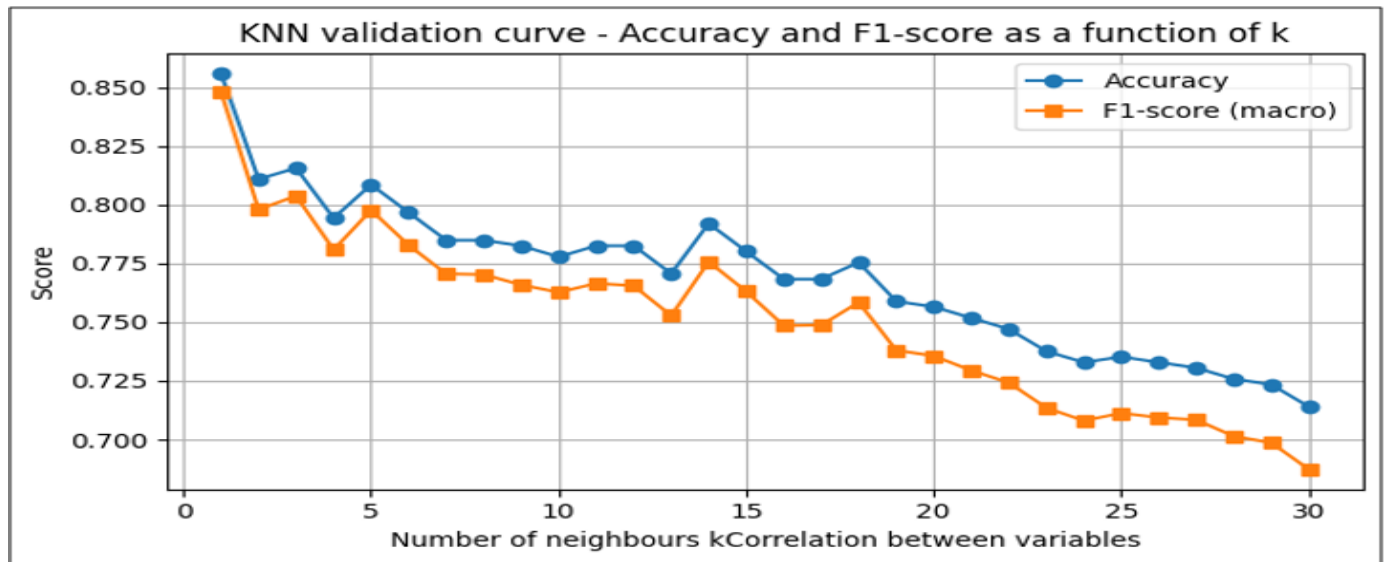


Fig 4 Validation curve Showing the Evolution of Accuracy and F1-score as a Function of the Number of Neighbors k ($k \in [1, 30]$)

Figure 4 illustrates the impact of the parameter k on the performance of the K-Nearest Neighbors classifier. It is observed that both accuracy and macro F1-score are maximized for $k=1$, reaching approximately 85.5% and 84.9%, respectively. Beyond this point, performance gradually declines, suggesting a loss of precision in class discrimination.

C. Feature Importance and Interpretability

The comparative analysis of feature importance across models represents a critical step in identifying consistent biomedical determinants and potential algorithmic biases.

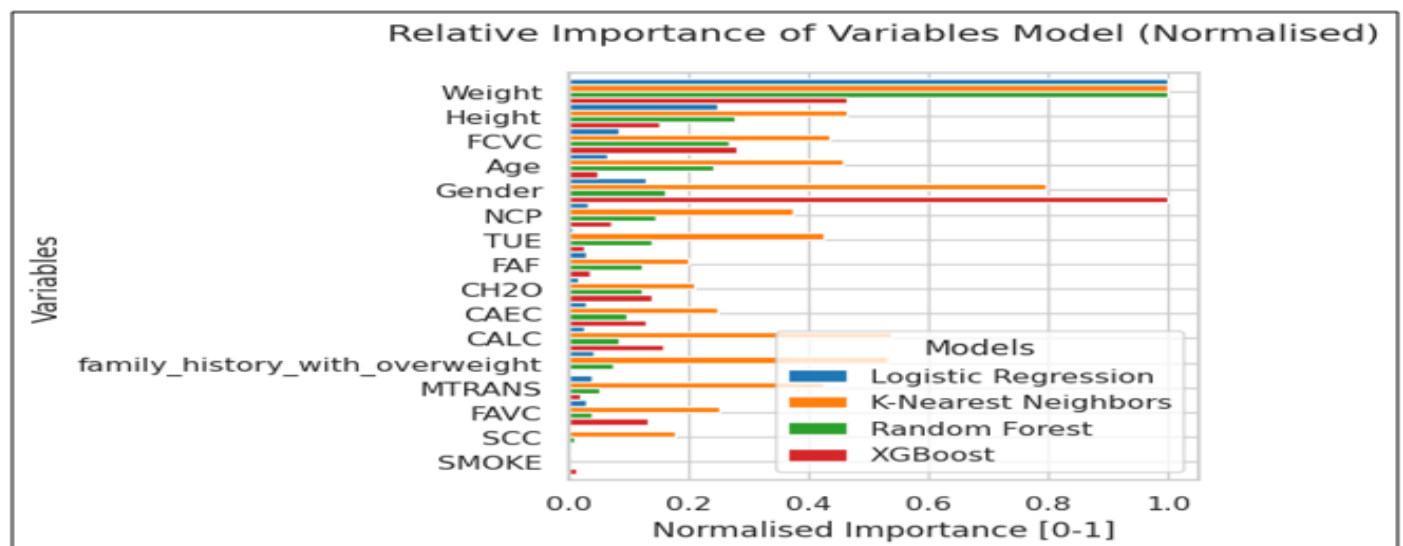


Figure 4: Relative Importance of Features Across Models

Figure 4 presents the normalized feature importance in our study, highlighting Weight, Height, and Frequency of Vegetable Consumption (FCVC) as the most influential variables across all models. However, notable differences are observed: XGBoost and Random Forest assign greater importance to Age and Gender, whereas Logistic Regression and K-Nearest Neighbors place more emphasis on factors such as Weight and Height.

D. Explainability Insights (SHAP Analysis)

➤ SHAP Summary Plots,

This section discusses the dominant variables identified through SHAP analysis, highlighting key factors such as Weight, Height, Age, and Frequency of Vegetable Consumption (FCVC) as the most influential in predicting obesity levels across the evaluated models.

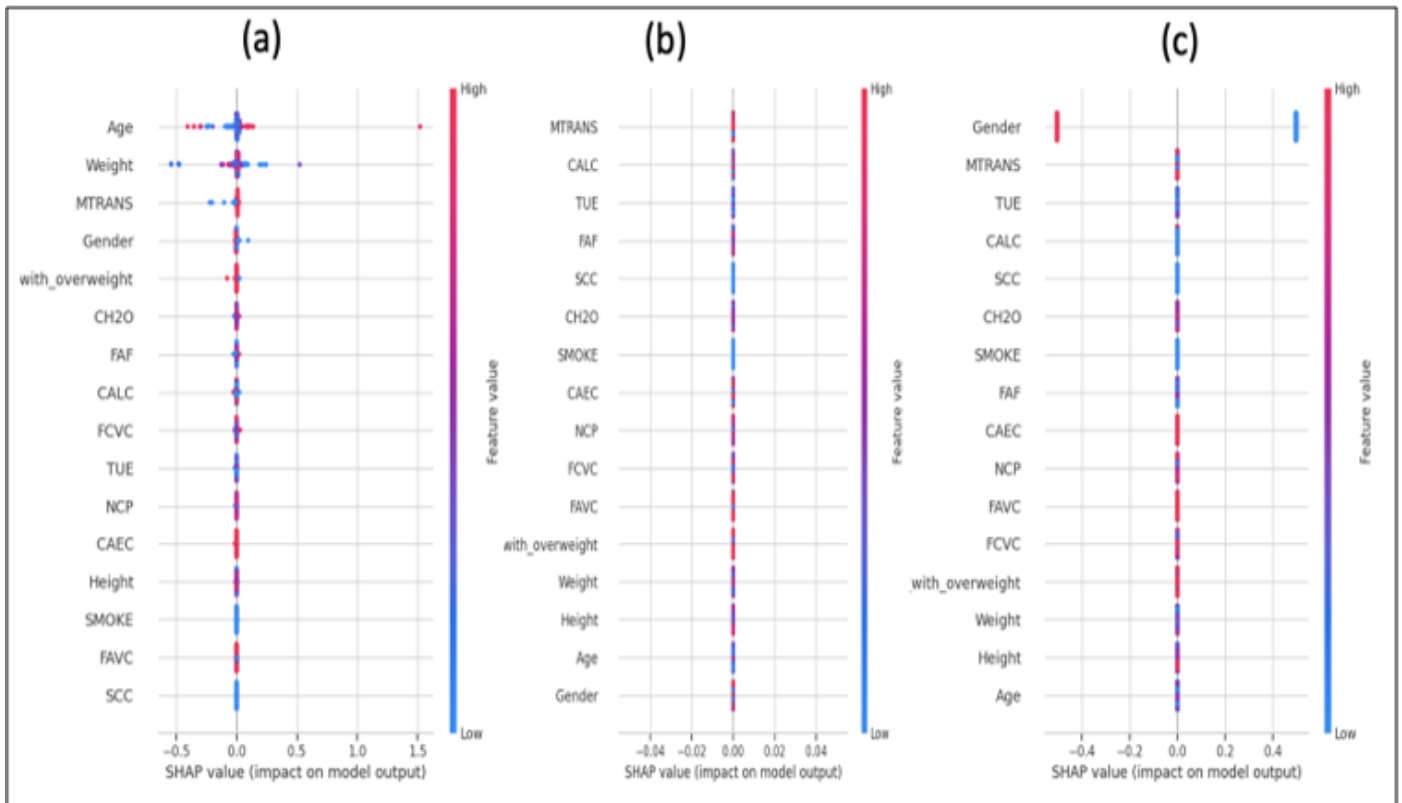


Fig 5 Relative Importance of Variables Across Models

Figure 5 illustrates the impact of variables on model predictions for K-Nearest Neighbors (a), Logistic Regression (b), and XGBoost (c) using SHAP values. The color gradient represents the feature values, with blue indicating low values and red indicating high values, providing a clear visualization of how each variable influences the classification outcomes.

- *SHAP Stacked Bar Charts for Multi-Class (MLP)*
A comparative analysis of feature influence across predicted classes, highlighting how each variable contributes differently depending on the obesity level classification.

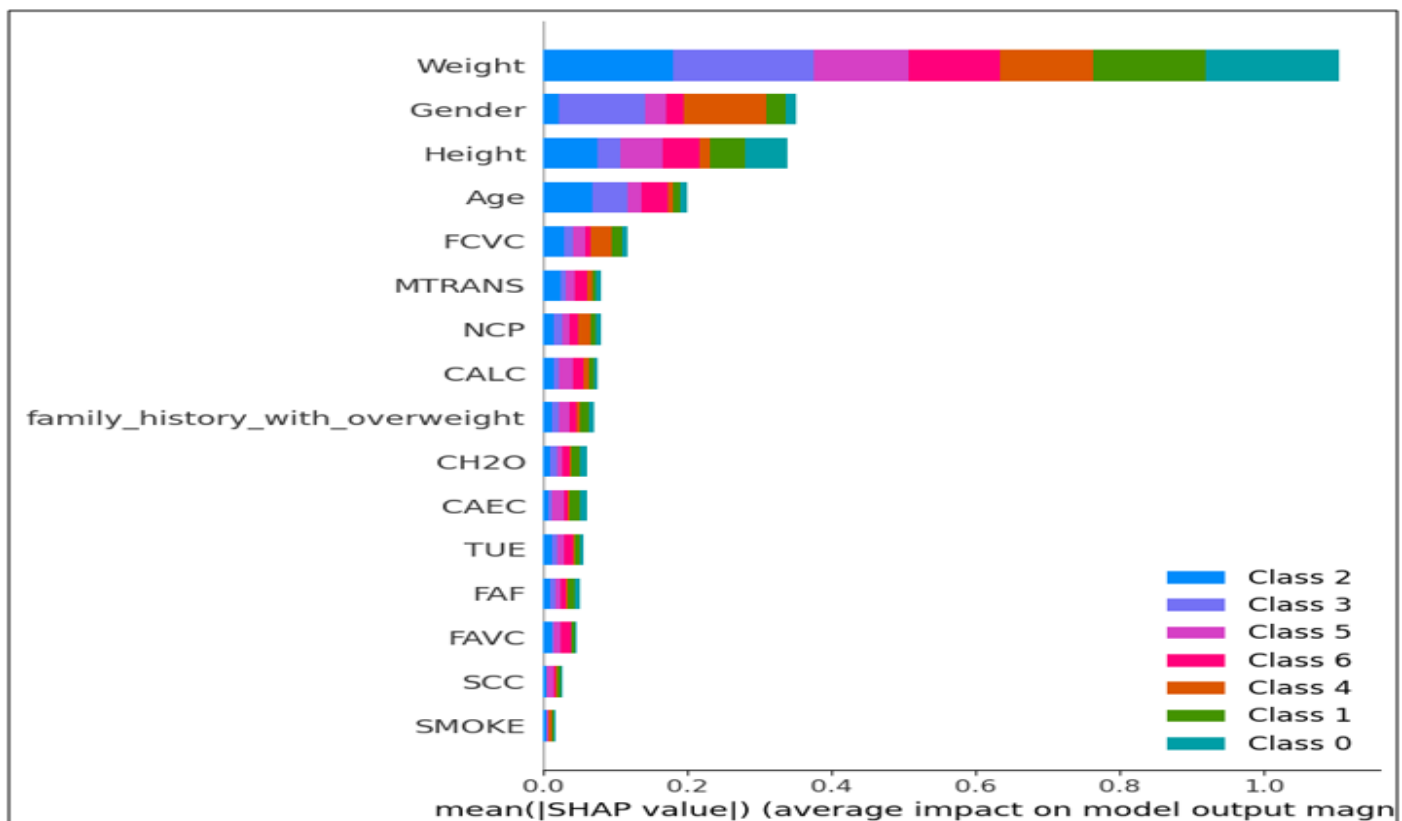


Figure 6: SHAP summary plots

Figure 6 visualizes the average feature importance for each target class (from Underweight to Obesity Type III). The chart highlights a clear dominance of certain behavioral and biometric variables across classifications.

V. DISCUSSION

The results of our study demonstrate that advanced models such as XGBoost and MLP outperform traditional approaches like Logistic Regression and KNN, achieving accuracies of 95.7% and 98.4%, respectively. These models effectively capture complex non-linear relationships, consistent with the findings of Sajid et al. (2022) and Esteve et al. (2019). However, KNN remains limited by the curse of dimensionality and class imbalance issues, as noted by Radovanović et al. (2009). While hyperparameter optimization through GridSearchCV significantly enhanced model performance, generalizing these results to other populations remains challenging (Brown et al., 2024). The integration of SHAP improved model interpretability by identifying key variables such as weight, height, and age, aligning with the recommendations of Lundberg and Lee (2017) and Samek et al. (2017). Nevertheless, SHAP's computational complexity on large datasets restricts its applicability in real-time clinical environments, as highlighted by Tjoa and Guan (2021).

VI. LIMITATIONS AND FUTURE RESEARCH

Although this study demonstrated the effectiveness of advanced models in predicting obesity risk, several limitations must be considered. The use of a single dataset from Kaggle may introduce selection bias, particularly due to limited sample diversity. Furthermore, while hyperparameter optimization improved model performance, the optimal configuration may not be generalizable to other populations. The integration of SHAP (SHapley Additive exPlanations) for model interpretability presents computational challenges, especially for complex models such as XGBoost and MLP, limiting its real-time application in clinical settings. Finally, the clinical acceptability of AI-driven decision support systems remains a major challenge, particularly for critical healthcare decisions.

Future research should focus on improving interpretability, incorporating temporal models such as LSTMs or Transformers, and expanding datasets to include greater demographic and geographic diversity. Additionally, integrating IoT (Internet of Things) sensors for real-time data collection could enhance prediction personalization.

VII. CONCLUSION

This study explores the use of artificial intelligence for obesity prevention, a rapidly growing global health concern. By comparing traditional models (logistic regression, KNN, Random Forest) with advanced approaches (XGBoost, MLP), we demonstrated that deep neural networks outperform classical methods in multiclass classification. All models were optimized using GridSearchCV, significantly improving their performance, with XGBoost achieving an accuracy of 96.9% and our optimized MLP reaching 98.05%.

Logistic regression, while useful, remains limited in modeling complex interactions, and KNN exhibits weaknesses in handling class imbalances. The integration of SHAP (SHapley Additive exPlanations) enhanced model interpretability by highlighting the importance of key variables such as weight, height, and age. Our hybrid approach, combining high performance, transparency, and generalizability, provides a robust and reproducible pipeline applicable in clinical settings.

REFERENCES

- [1]. World Health Organization. (2021). Obesity and overweight. <https://www.who.int/news-room/fact-sheets/detail/obesity-and-overweight>
- [2]. Guh, D. P., et al. (2009). The incidence of comorbidities related to obesity and overweight: A systematic review and meta-analysis. *BMC Public Health*, 9, 88. <https://doi.org/10.1186/1471-2458-9-88>
- [3]. Maria, A.S & Ramasamy, Sunder & Kumar, R.Satheesh. (2023). Obesity Risk Prediction Using Machine Learning Approach. 1-7. 10.1109/ICNWC57852.2023.10127434.
- [4]. Lin W, Shi S, Huang H, Wen J, Chen G. Predicting risk of obesity in overweight adults using interpretable machine learning algorithms. *Front Endocrinol (Lausanne)*. 2023 Nov 17;14:1292167. doi: 10.3389/fendo.2023.1292167. PMID: 38047114; PMCID: PMC10693451.
- [5]. M. Dirik, "Application of machine learning techniques for obesity prediction: a comparative study," *Journal of Complexity in Health Sciences*, Vol. 6, No. 2, pp. 16–34, Oct. 2023, <https://doi.org/10.21595/chs.2023.23193>
- [6]. Lundberg, Scott & Lee, Su-In. (2017). A Unified Approach to Interpreting Model Predictions. 10.48550/arXiv.1705.07874.
- [7]. Saxena, A.; Mathur, N.; Pathak, P.; Tiwari, P.; Mathur, S.K. Machine Learning Model Based on Insulin Resistance Metagenes Underpins Genetic Basis of Type 2 Diabetes. *Biomolecules* 2023, 13, 432. <https://doi.org/10.3390/biom13030432>
- [8]. Sadiku, Matthew & Eze, Kelechi & Musa, Sarhan. (2018). Data Mining in Healthcare. *International Journal of Advances in Scientific Research and Engineering*. 4. 90-92. 10.31695/IJASRE.2018.32881.
- [9]. E. Tjoa and C. Guan, "A Survey on Explainable Artificial Intelligence (XAI): Toward Medical XAI," in *IEEE Transactions on Neural Networks and Learning Systems*, vol. 32, no. 11, pp. 4793–4813, Nov. 2021, doi: 10.1109/TNNLS.2020.3027314.
- [10]. Mehrparvar, F. (2023). Obesity Level Dataset. Kaggle. <https://www.kaggle.com/datasets/fatemehmehrpavar/obesity-levels>.
- [11]. Begum, N. B. M., et al. (2024). Predicting Obesity Levels: Machine Learning Approaches Analyzing Eating Habits and Lifestyle Factors. *NeuroQuantology*, 22(4), 305-313. <https://www.neuroquantology.com/open-access/PREDICTING%2BOBESITY%2BLEVELS.DOI> Number: 10.48047/nq.2024.22.4.nq24035.

- [12]. Alsareii, Saeed & Awais, Muhammad & Alamri, Abdulrahman & Alasmari, Mansour & Irfan, Muhammad & Raza, Mohsin & Manzoor, Umer. (2023). Machine-Learning-Enabled Obesity Level Prediction Through Electronic Health Records. *Computer Systems Science and Engineering*. 46. 3715-3728. 10.32604/csse.2023.035687.
- [13]. Helforouh Z, Sayyad H. Prediction and classification of obesity risk based on a hybrid metaheuristic machine learning approach. *Front Big Data*. 2024 Sep 30;7:1469981. doi: 10.3389/fdata.2024.1469981. PMID: 39403430; PMCID: PMC11471553.
- [14]. Genc, A. C., & Arican, E. (2025). Obesity classification: a comparative study of machine learning models excluding weight and height data. *Revista da Associação Médica Brasileira*, 71(1), e20241282. <https://doi.org/10.1590/1806-9282.20241282>.
- [15]. Hosmer Jr, D. W., Lemeshow, S., & Sturdivant, R. X. (2013). *Applied logistic regression*. John Wiley & Sons.
- [16]. Breiman, L. (2001). Random forests. *Machine Learning*, 45(1), 5-32. <https://doi.org/10.1023/A:1010933404324>
- [17]. Chen, T., & Guestrin, C. (2016). XGBoost: A Scalable Tree Boosting System. *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, 785–794. <https://doi.org/10.1145/2939672.2939785>.
- [18]. Radovanovic, Milos & Nanopoulos, Alexandros & Ivanovic, Mirjana. (2009). Nearest neighbors in high-dimensional data: The emergence and influence of hubs. *Proceedings of the 26th International Conference On Machine Learning, ICML 2009*. 382. 109. 10.1145/1553374.1553485.
- [19]. Chaoyu Gong, Zhi-gang Su, Xinyi Zhang, Yang You, Adaptive evidential K-NN classification: Integrating neighborhood search and feature weighting, *Information Sciences*, Volume 648, 2023, 119620, ISSN 0020-0255, <https://doi.org/10.1016/j.ins.2023.119620>.
- [20]. Zhang, Chao & Zhong, Peisi & Liu, Mei & Song, Qingjun & Liang, Zhongyuan & Wang, Xiao. (2022). Hybrid Metric K-Nearest Neighbor Algorithm and Applications. *Mathematical Problems in Engineering*. 2022. 1-15. 10.1155/2022/8212546.
- [21]. Nitish Srivastava, Geoffrey Hinton, Alex Krizhevsky, Ilya Sutskever, and Ruslan Salakhutdinov. 2014. Dropout: a simple way to prevent neural networks from overfitting. *J. Mach. Learn. Res.* 15, 1 (January 2014), 1929–1958.
- [22]. Sergey Ioffe and Christian Szegedy. 2015. Batch normalization: accelerating deep network training by reducing internal covariate shift. In *Proceedings of the 32nd International Conference on International Conference on Machine Learning - Volume 37 (ICML'15)*. JMLR.org, 448–456.
- [23]. Raghu, Maithra & Zhang, Chiyuan & Kleinberg, Jon & Bengio, Samy. (2019). Transfusion: Understanding Transfer Learning with Applications to Medical Imaging. 10.48550/arXiv.1902.07208.
- [24]. Kingma, Diederik & Ba, Jimmy. (2014). Adam: A Method for Stochastic Optimization. *International Conference on Learning Representations*.
- [25]. Prechelt, Lutz. (2000). Early Stopping - But When?. *Lecture Notes in Computer Science*. 10.1007/3-540-49430-8_3.
- [26]. M. Sajid et al., Deep learning approaches for predicting obesity using dietary and physical activity patterns. *Computers in Biology and Medicine*, vol. 149, 105962, 2022. DOI: 10.48047/nq.2024.22.4.nq24035.
- [27]. Esteva A, Robicquet A, Ramsundar B, Kuleshov V, DePristo M, Chou K, Cui C, Corrado G, Thrun S, Dean J. A guide to deep learning in healthcare. *Nat Med*. 2019 Jan;25(1):24-29. doi: 10.1038/s41591-018-0316-z. Epub 2019 Jan 7. PMID: 30617335.
- [28]. Brown, Williams & Noah, Asher & John, Ada. (2024). Optimizing Hyperparameters in Machine Learning Models: Techniques and Best Practices.
- [29]. Samek, Wojciech & Wiegand, Thomas & Müller, Klaus-Robert. (2017). Explainable Artificial Intelligence: Understanding, Visualizing and Interpreting Deep Learning Models. *ITU Journal: ICT Discoveries - Special Issue 1 - The Impact of Artificial Intelligence (AI) on Communication Networks and Services*. 1. 1-10. 10.48550/arXiv.1708.08296.