

EXPLAINABLE MACHINE LEARNING FOR EVALUATING DIABETES PREDICTION MODELS

Saumyaa Dhakan*

12th Grade, IB Diploma Programme, Fountainhead School, Surat, India

*Corresponding author:

saumyaadhakan9@gmail.com

Abstract

The present study has investigated the use of Explainable AI (XAI) machine learning models for diabetes prediction, with a particular focus on SHAP (Shapley Additive Explanations) values and XGBoost importances (gain). Using the publicly available Pima Indians Diabetes dataset there were various classification models: Logistic Regression, Random Forest, and XGBoost trained and evaluated based on the parameters: Accuracy, and AUC. Global and local explainability was further assessed through SHAP bar plots and XGboost gain importances plots were used for analyzing structural feature importances. Results gathered suggested that stabilized glucose level and age were the most influential predictors of diabetes across models even though differences emerged in the ranking of the feature HDL (High-Density Lipoprotein) cholesterol. Overall, while SHAP provided a hybrid understanding of local and global patient-specific explanations, XGboost gain importances gave structural importance. The discussion section also includes how the combination of these two tools provided a hybrid, comprehensive and nuanced understanding about how these models bridge the “black-box” gap in clinical decision-making, ensuring trust, clarity and transparency in the prediction of AI models for sensitive domains like healthcare applications.

Keywords: Diabetes prediction, Explainable AI (XAI), SHAP, XGBoost importances, machine learning, interpretability, healthcare, classification models.

I INTRODUCTION

Machine learning is used in various spheres like healthcare to support disease prediction, early diagnosis and personalized treatment. It has been used significantly more nowadays worldwide. Diabetes is a commonly known chronic disease where a person experiences high blood sugar levels because the pancreases are not able to produce enough insulin or the body's inability to use the insulin effectively. It is a metabolic disease which is growing in prevalence worldwide, predictive models play a major role in risk-assessment and timely intervention and measures. Early and accurate prediction of diabetes can help clinicians identify at-risk patients and guide preventive measures and improve long-term patient outcomes.

However, most of the models function as “black boxes” where they just provide the prediction without providing insight of the specific factors that contributed to outcomes. In clinical practice, it lacks transparency posing significant challenges. Physicians don't just need to know whether a patient is at risk but why the model assigned the risk due to which factor. For example, biomarkers like stabilized glucose level, HDL, cholesterol, body size may have well-established clinical relevance, but their influence on prediction may vary across different models. Therefore, without interpretability clinicians may find it difficult to trust and adopt these AI- driven tools and predictive models.

This is where Explainable AI (XAI) becomes necessary in bridging this gap. By using tools such as SHAP(Shapely Additive explanations) alongside model-specific measures like XGboost feature importance plot, this research paper aims to balance predictive accuracy with interpretability.

This research uses the publicly available [1] Pima Indians Diabetes Dataset by National Institute of Diabetes and Digestive and Kidney Diseases (NIDDK) and distributed by UC Irvine's Machine Learning Repository and apply different ML classification models such as Logistic Regression, Random Forest, and XGBoost. The models are compared based on the parameters:accuracy and AUC, followed by explainability analysis through SHAP values to visualize and understand the contribution of features such as age, marital status, loan, and campaign interactions.

The motivation of this paper comes from the necessity for predictive models that are not only precise but also comprehensible and reliable within healthcare settings. This paper aims to demonstrate responsible integration of AI to clinical decision-making by focusing on diabetes prediction and thoroughly evaluating various explainability methods, to make sure the model strengthens rather than dispels medical expertise.

The flow of this paper is as follows: Section 2 presents the related work and background on explainable AI. Section 3 talks about the dataset and the method, including preprocessing and training the model. Section 4 talks about the results of comparing models and looking at how explainable they are. Section 5 discusses the research's and potential research's limitations. Finally, Section 6 is the conclusion that highlights the significance explainable AI plays in making transparent and trustworthy decision making systems.

II Related work

A. Kaliappan et al. (2024) – “Analyzing classification and feature selection strategies for diabetes prediction”

In this study [3] model explainability tools: Random Forest, XGBoost, SVM; global SHAP bar plots and XGBoost-gain importance were used to analyze multiple public diabetes dataset(including the Pima dataset used in this research paper). The research paper also provides basic workflow of pre-processing from Xgboost to SHAP bar which is evident in the inclusion of sections, “Data processing”, “Model training with Xgboost” and “Explainability with SHAP.” This workflow allows the readers to have a comprehensive understanding about the differences in models in SHAP bar plots presented in this study making it similar to this study where different models of SHAP bar plots are created for understanding feature ranking. The advantages of this study include that the use of multiple datasets remove the margin of error of any bias or misleading conclusion providing better idea of how the models provide different SHAP bar plots with different feature ranking for a common prediction i.e diabetes. However, this study is focused only on comparing feature ranking of each model's SHAP bar plots but does not take into account local SHAP explanations i.e. the magnitude of the SHAP value(feature's influence) in the prediction. Alongside there is no inclusion of comparison of quantitative comparison of SHAP vs. gain values for nuanced understanding of each feature's influence in prediction of diabetes in different models and datasets.

B. Kibria et al. (2022) – “An Ensemble Approach for the Prediction of Diabetes Mellitus Using a Soft-Voting Classifier with an Explainable AI”

In this study [4] six classifiers: ANN, RF, SVM, LR, AdaBoost, XGBoost are combined in soft-voting ensembles using the Indiana diabetes datasets used in this research paper as well. The research paper demonstrates that global explainability of SHAP can be retained and understood from ensembles as they are predicting probabilities averaged from various models evident in the “Performance Analysis and Experimental Results” section of the study. The advantages of this study include strong performance of the ensemble prediction with accuracy of 90% giving the readers a clear global SHAP explanation for diabetes predictions. However, the study does not take into account local explainability things like SHAP rankings with XGBoost-gain importance that are important for a comprehensive understanding of the predictions and features involved for the same.

C. Liu et al. (2022) – “Predicting the Risk of Incident Type 2 Diabetes Mellitus in Chinese Elderly Using Machine Learning Technique”

In this study [5] models: Logistic Regression, Decision Tree, Random Forest, XGBoost are used to predict the top predictors of Type 2 Diabetes by using 127 031 records of Chinese seniors (BRFSS-derived). Through the use of these models this research paper highlights the importance of SHAP in identifying how SHAP accurately identifies clinically

meaningful biomarkers (glucose, waist) similar to this research paper study. The advantages of this study include its large demographic diverse inclusion of data helping in reducing marginal errors, possibilities of bias or misleading conclusions. However, it does not take into account XGBoost-gain plot importances graphs and has a limited focus on SHAP values by just analyzing the global importance without local case studies.

Overall, all the three studies provide the inspiration for analyzing SHAP and Xgboost importances for the Pima dataset to understand the predictors that actually influence the risk of diabetes. By reading each study one thing is clear i.e to focus on various aspects of SHAP which is analyzing the global as well local interpretability values of the features in the dataset alongside comparing it with the XGBoost-gain (tree-split contribution). This would help achieve the aim of the study i.e to understand how explainable AI thorough global and local explanations bridge the “black-box” gap for clinicians.

III IMPLEMENTATION

A. Dataset

The dataset [1] is related to clinical and demographic records of patients for diabetes prediction. It was collected to study medical indicators such as glucose, cholesterol, and HbA1c levels, along with age and gender, in order to evaluate their role in assessing diabetes risk.

Table 1 shows different attributes of the dataset.

Attribute	Meaning
chol	Total cholesterol level in the blood (mg/dL).
hdl	High-Density Lipoprotein (HDL) cholesterol – often called "good" cholesterol.
stab.glu	Stabilized (fasting) blood glucose level – blood sugar after fasting, measured in mg/dL
ratio	Cholesterol-to-HDL ratio – used to assess cardiovascular risk.
glyhb	Glycated hemoglobin (HbA1c) – indicates average blood glucose over the past 2-3 months. Note: this attribute highly affects the output target
location	Geographic location of the patient (e.g., Buckeye, Akron, Cleveland)
age	Age of the patient
gender	Biological sex of the patient (male/female)
height	Height of the patient in inches
weight	Weight of the patient in pounds
frame	Body frame size – typically small, medium, or large
bp.1s	First systolic blood pressure reading – pressure when the heart beats (mm Hg).
bp.1d	First diastolic reading – pressure between heartbeats (mm Hg)
bp.2s	Second systolic reading – a follow-up or second visit reading.
waist	Waist circumference – important for assessing central obesity (inches).
hip	Hip circumference – used along with waist to calculate Waist-to-Hip Ratio (WHR), a health indicator.
time.ppn	Time (in minutes) since the patient last saw a healthcare provider, reported a symptom, or had a medical consultation.

Table 1

B. Correlation matrix

The correlation matrix shows the relationship between different independent variables in the dataset using correlation coefficients. The values range from -1 to +1, where values closer to +1 indicate a strong positive correlation and values closer to -1 indicate strong negative correlation. The diagonal is always equal to 1 since each variable perfectly correlates to itself. This matrix is structured with rows and columns allowing for easy comparison of data. These rows and columns depict the correlation value of each variable. It is a powerful tool to summarize large datasets, identify hidden patterns, and visualize relationships that can influence outcomes (diabetes).

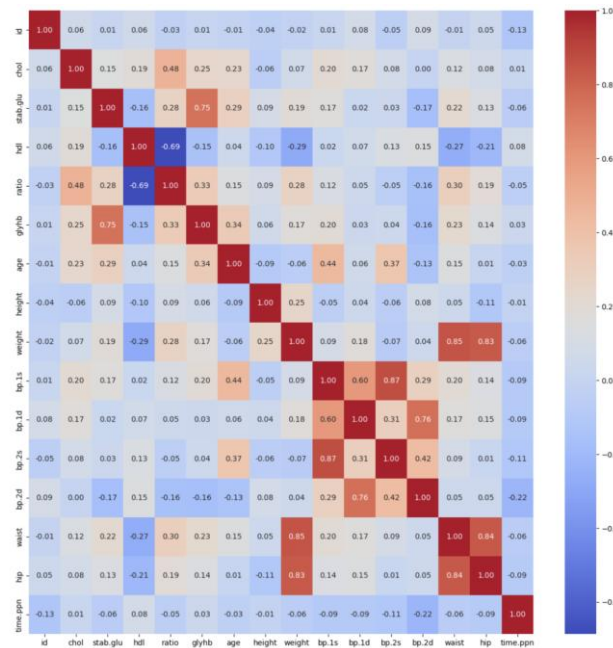


Fig 1. Correlation matrix

Based on Fig 1, it can be observed that the variables: weight and waist & weight and hip & weight and hip show very high correlation values (above 0.8). Similarly the variables: bp.1s and bp.2s also display a strong positive correlation. This indicates that these features are closely related and tend to vary collectively, reflecting some psychological link such as body composition and blood pressure measurement due to these variables.

C. Models

To develop a reliable prediction system for diabetes there is implementation of a range of machine learning models. Each model has its own strengths in terms of interpretability, complexity and prediction performance. First, is Logistic Regression which estimates the probability of diabetes based on input features like glucose, insulin, BMI, and waist circumference making it simple and accurate in its interpretability. Second, is Decision tree that works on splitting the dataset into branches based on feature thresholds (e.g., glucose > 5.7) until it reaches a decision; however, it has risks of overfitting. Further, is Random Forest classifier, which is an ensemble of Decision Trees where each tree is trained on random subset of data and features improving accuracy and reducing risk of overfitting. Naïve Bayes is another model that works on the Bayes' theorem with the assumption that the features are independent of each other. K-Nearest Neighbors (KNN) is a distance based model that works on predicting new patients based on the majority class of the most similar patients in the dataset. Gradient Boost classifier is based on creating decision trees in a sequential manner where each reduces the margin of error made by the previous ones. Last but not the least is Extreme Gradient Boosting(XGBoost), which is an optimized implementation of Gradient boost designed for more efficiency and performance. Its power lies in prediction, and finding missing data effectively.

All the mentioned models when combined would provide interpretable feature contributions. In the upcoming section, the performance of these models will be evaluated using Accuracy and Area Under Curve (AUC). Based on the results the two-best performing models will be identified and interpreted using SHAP bar plots to understand the contribution of individual features to prediction of diabetes.

D. Parameters for Gradient boost

Parameters	Value
learning_rate	0.1
n_estimators	100
Max_depth	3
min_sample_split	2

These are the standard default parameters of the Gradient boost that contribute to the model's performance, trade-off between biases and variances in the dataset and risk of overfitting.

Learning rate regulates the contribution of each new decision tree(diabetic vs non-diabetic groups) to the ensemble allowing the model to improve diabetes prediction gradually & avoiding overfitting. N_estimators are the total boosting stages of the model that are the number of decision trees combined sequentially to refine the diabetes prediction. Min_sample_splits is the minimum number of patient records required in the decision tree to split an internal node. Max_depth is the maximum number of levels of splits the tree can make.

E. Parameters for XGboost

Parameters	Value
learning_rate	0.1
Gamma	0
Max_depth	6
min_child_weight	1

These are the default parameters of the XGboost that contribute to the model's performance, trade-off between biases and variances in the dataset and risk of overfitting.

Learning rate regulates the contribution of each new decision tree(diabetic vs non-diabetic groups) to the ensemble allowing the model to improve diabetes prediction gradually & avoiding overfitting. Gamma is the minimum reduction in loss required for making a split; this prevents the model from creating unnecessary branches on weak predictors (skin thickness or pregnancies) & instead emphasizes on stronger ones (like glucose levels). Max_depth is the number of splits each decision tree can make. Higher max_depth value allows for better analysis of complex patterns. Min_child_weight is the minimum sum of instance weights(patient records) required in a child node of decision trees.

IV QUANTITATIVE ANALYSIS:

This analysis explores comparing values of accuracy and AUC of all the models. Table 1 display accuracy and AUC are the main metrics used to evaluate the performance. Accuracy measures the proportion of correctly classified predictions and AUC evaluates the ratio of correctly identified diabetic as diabetic(True positive rate) and incorrectly identified non-diabetic patients as diabetic(False positive rate).

Model	Accuracy	AUC
Random Forest	0.8926	0.9265
Logistic Regression	0.9008	0.9488
XGboost	0.8760	0.8976
Naive Bayes	0.8760	0.9204
Decision Tree	0.8347	0.8276
KNN	0.8678	0.8329
Gradient boost	0.8926	0.8948

Table 1

Table 1 shows that Logistic regression generated the highest accuracy score and AUC. To further understand and compare the model it can be visualized with the AUC graph.

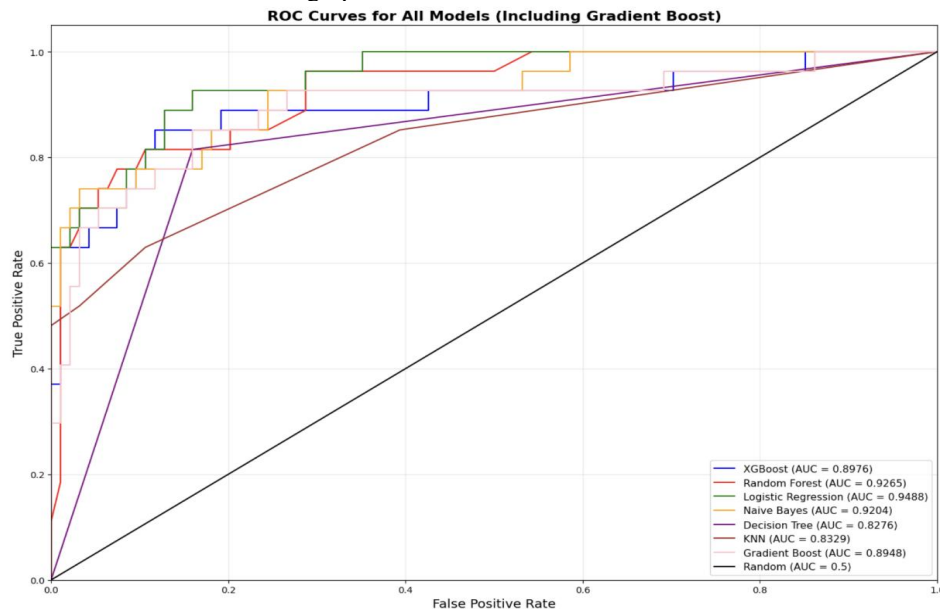


Fig 2: AUC score of models

Figure 2 showcases AUC curves and values for each model. From this it can be clearly observed that Logistic Regression has the highest AUC curve (green color). Higher AUC values allows better performance of the model as it can correctly classify more diabetics (True Positives) while avoiding false alarms (False Positives).

V QUALITATIVE ANALYSIS:

In a machine learning model, understanding each variable's influence and importance in the dataset prediction result is essential.

SHAP stands for Shapely Additive explanations. It is a tool working on game theory that helps us quantify each dataset's feature contribution to the prediction which in this case is whether the patient has diabetes(yes or no). It can be visualized with the help of SHAP bar plot. This plot allows us to understand global interpretations i.e to understand feature's effect across the entire dataset.

Furthermore, it is essential to analyze & discuss SHAP models (bar plots) of the top 2 models for further comparison which are Logistic Regression & Random Forest.

A. SHAP Bar plot of Logistic Regression

Below is the SHAP barplot of Logistic Regression, our best-performing model.

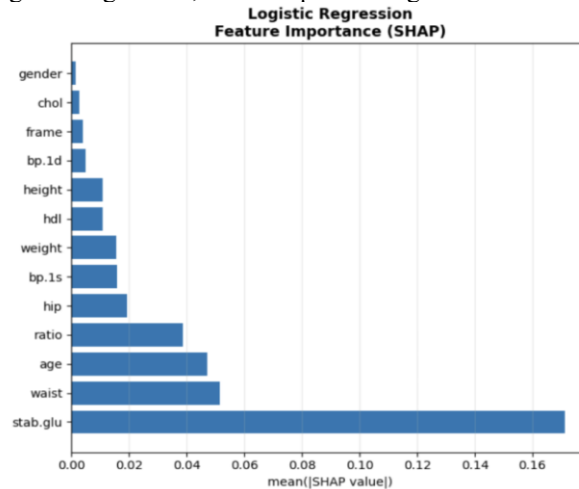


Fig 3: SHAP bar plot of Logistic Regression.

From Fig 3, it can be interpreted that the feature stab.glu i.e glucose level of the patient is the highest feature, because it reflects inefficiency of insulin regulation making it the most impactful, realizable predictor for predicting diabetes. If a person has high glucose level the chances for the patient having diabetes is high as there is weak insulin-regulation. The second-most impactful feature is waist size/circumference of the patient. Higher waist size indicates excess abdominal fat that increases insulin resistance making it more likely that the patient has diabetes.

B. SHAP Bar plot of Random Forest

Below is the SHAP barplot of Random Forest, the second best-performing model.

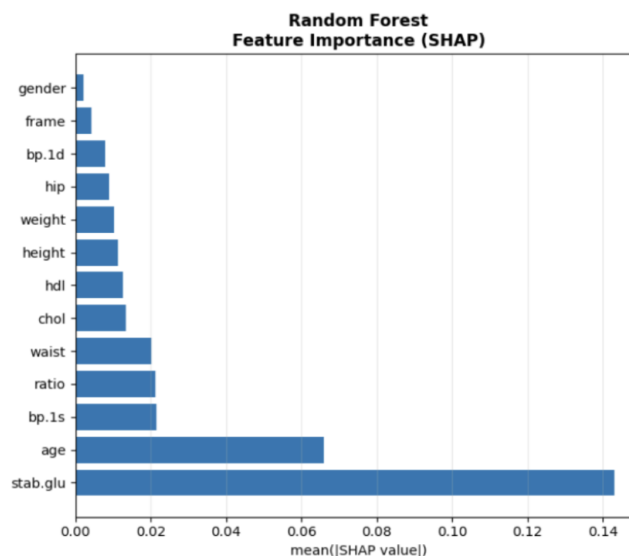


Fig 3: SHAP bar plot of Random Forest.

From Fig 3, it can be interpreted that the feature stab.glu i.e glucose level of the patient is the highest feature as similarly observed in the Logistic Regression model. If a person has high glucose level the chances for the patient having diabetes is high as there is weak insulin-regulation. The second-most impactful feature is varied in this model which is age of the patient. Age is another relevant and important feature that plays a major role in the prediction of diabetes. As a person ages the risk for specifically Type 2 diabetes increases due to a decline in the insulin-sensitivity (ability of the body to respond to insulin) increasing the glucose content in patients blood. Therefore, with age the chances of being diabetic increase.

C. XGBoost SHAP bar plot

After SHAP, the most important model in explainability of AI is XGboost. Figure 3, visualizes XGboost feature importance graph specifically type 'gain.' Type gain graphs, helps in understanding the contribution of each feature in reducing prediction error. From the graph, it can be seen that the feature 'stab.glu' i.e glucose level of the patient has the most contribution in reducing marginal errors in prediction of whether a patient has diabetes or not. The second-most important feature is the age of the patient. Age of the patient determines the risk they have to be diagnosed with Type 2 diabetes due to decrease in insulin sensitivity as we age.

XGBoost - Top 15 Feature Importances (Gain)

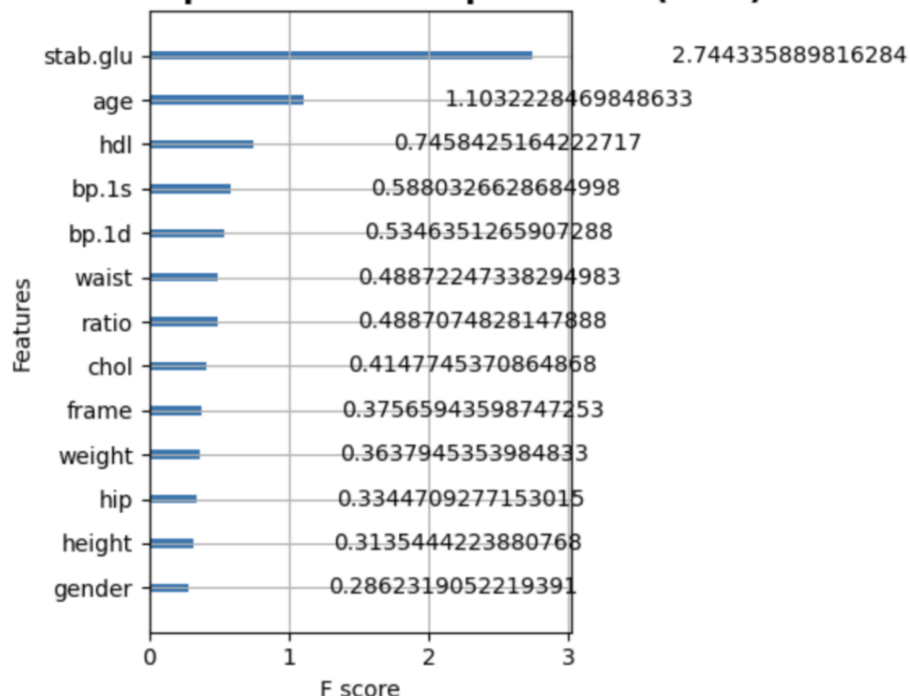


Fig 3: XGboost feature importances (gain)

V. Discussion

In this section, the analysis of the final results are summarized, alongside a comparison between the results obtained in SHAP plots and XGboost plots with context of explainable AI, highlighting strengths and weaknesses of each approach used. Finally, key takeaways and limitations are presented.

The SHAP bar plot analysis for the Logistic Regression model indicates that stab.glu(stabilized glucose level) is the strongest predictor for diabetes, which is followed by waist-size, age and body fat ratio, on the other hand features like gender, cholesterol, and frame contribute significantly lower. Overall, the model concludes that metabolic and anthropometric (measurements of human-body) factors are more influential over demographic ones. Moreover, the SHAP bar plot of the Random forest model suggests that stab.glu(stabilized glucose level) is the most important predictor of diabetes, followed by age, with moderate contributions from blood pressure (bp.1s, bp.1d), waist and ratio. Other features such as gender, frame and weight have limited influence on the diabetes prediction. Results from the XGboost feature importance gain graph highlights that biological markers (glucose, lipids, blood pressure) and demographic factors (like age) significantly influence the prediction of diabetes. Other factors such as lifestyle or structural measures (like height, weight, gender) are weak predictors in the dataset meaning their influence on the prediction is minimal.

The following discussion would talk about comparing SHAP values with XGboost importances to better understand the explainability of models. Both SHAP analysis and XGBoost feature importance prove the trend of stabilized glucose (stab.glu) as the most dominant predictor of diabetes which is followed by age and HDL cholesterol and other features such as blood pressure, waist size, body -fat ratio contributing moderately. This commonality suggests that these features(like stabilized glucose level, age)are not only structurally important for the model's decision-making (as shown by gain) but also have consistent predictive influence across individual instances (as captured by SHAP) of prediction of diabetes.

However, the feature HDL(High-Density Lipoprotein (HDL) cholesterol) was ranked differently in both the importances. As, SHAP offers a global view though magnitude value of all averaged patients (in the form of bar plots) & local interpretability by breaking down contributions for a single patient's prediction both the models (Random forest & Logistic Regression) ranked this feature lower. This indicates that its direct contribution to the prediction of diabetes across the dataset was relatively limited. However, SHAP's local explanations suggest that for certain individual patients, HDL could still have a noticeable effect on their predicted risk of diabetes diagnosis. In contrast, XGBoost feature importance (gain) analysis, HDL was ranked notably higher indicating a difference in how both the models are unique in

methodology in measuring the contribution feature importances and explainability. This discrepancy is relevant in clinical histories as HDL is a well-known biomarker in one's metabolic health, and its importance shown in one method but not in other shows the need to interpret model output cautiously in a healthcare context where each biomarker has a key role in patient's health.

While, XGboost's feature importance gain model measures how much each variable improves tree splits, SHAP values show a transition ahead by providing magnitude value of each feature's contribution to predictions higher or lower providing more interpretable explanation with local and global interpretability. This contrast between SHAP and XGboost feature importances highlights a very important difference between two models of explainability of AI: XGboost feature importance (gain) measures how a feature helps the model structure itself, whereas SHAP reflects the real impact of feature on final predictions. This explains that reliance on a single tool may lead to discrepancies giving an incomplete or misleading view. Applying both these methods allows having a comprehensive, hybrid and nuanced understanding of both global interpretability(overall) as well as local interpretability(specific to a patient) which is really essential when applying AI models in sensitive domains like clinical-testing.

Overall, this research paper explores the role of explainable AI (XAI) in modeling prediction of diabetes through analyzing feature contributions. Through both quantitative analysis (Accuracy and AUC) and qualitative analysis (SHAP and XGboost gain plots) it was observed that stab.glu(stabilized glucose) and age were the parameters that had the most impact on the prediction of whether a patient has diabetes or not across all the models. While SHAP importances highlighted the direct influence of features on the prediction, XGboost feature importances provided a structural view of how variables contribute to building the model.

This comparison revealed several relevant similarities as well as discrepancies. Evidence in case of HDL ranking which was projected lower in SHAP cross Logistic Regression and Random Forest suggesting a weaker influence on the prediction while XGboost feature importances ranking it higher showing its structural role in improving prediction errors. This accentuates that relying on a single explainability method may result in incomplete or misleading interpretations especially in the medical domain where biomarkers such as HDL have a clinical significance in a patient's health.

Key takeaways from the research suggest that to identify whether a patient is diagnosed with diabetes it is important to rely on the patient's stabilized glucose level and age supported by all the models analyzed across. The value of combining both the explainability models: SHAP and Xgboost allows for having a dual understanding of local and global interpretability and better prediction about diabetes. It is important to keep in mind clinical cautions since observed discrepancies in HDI importance in different models may also not align with the biomedical knowledge or patient scenario due to histories. Therefore, it is important to have careful interpretation by domain experts about a patient's diagnosis. Applying multiply interpretability tools together provides prediction models that are trustable, transparent and bridges the gap between model performance and clinical decision-making.

However, there are some limitations that need to be checked. Firstly, the analysis was based on only one dataset without external validation causing uncertainty about the generalizability of the findings to broader populations. Lack of validation of independent dataset reduces the ability to confirm robustness and stability of the identified important predictors. Moreover, the analysis was limited with traditional machine learning models (XGboost, Random Forest, and Logistic Regression). Although these models are computationally efficient, interpretable, they may not capture non-linear patterns that advanced deep learning approaches potentially could. Therefore, this limits the scope of the analysis and also limits the predictive accuracy and scalability of findings.

Thirdly, there are methodological constraints when it comes to the results obtained from SHAP feature importances. This is because SHAP explanations vary with model choice alongside feature importance metrics from XGBoost may also influence the values of variables that are correlated with each other. Finally, the analysis did not take into account other relevant data features of the patient such as longitudinal patient histories, genetic markers, or lifestyle information. By not including these features there may be restrictions in the overall scope of prediction model and clinical applicability.

VI. Conclusion

To conclude, the integration of SHAP and Xgboost feature importances provides a comprehensive, integrative view of machine learning models and features that are influential in prediction of diabetes. This can be expanded in the future by including datasets which are larger and have diverse scenarios for reducing statistical marginal errors. Alongside applying some more advanced approaches like LIME(Local Interpretable Model-Agnostic Explanations) an approach that helps us understand the decision mechanism of the black box models and other more advanced models. There can be engagement with clinical experts for a better evaluation process to make sure that the AI systems align with medical reasoning and patient safety.

ACKNOWLEDGEMENT

The research project was undertaken independently in pursuit of higher education. It holds no connection with the school and nor is it part of the school curriculum. The authors wish to acknowledge that there is use of CHatGPT for writing

research. It was utilised as a tool for assisting in improving the language and formatting of the paper. This paper remains as an accurate representation of the author's underlying work and novel intellectual contribution.

VI REFERENCES:

- [1] "Diabetes Dataset," *Kaggle*, IMT Kaggle Team. [Online]. Available: <https://www.kaggle.com/datasets/imtkaggleteam/diabetes>
- [2] E. K. Marsh, "Calculating XGBoost Feature Importance," *Medium*, Jan. 31, 2023. [Online]. Available: <https://medium.com/@emilykmarsh/xgboost-feature-importance-233ee27c33a4>
- [3] S. Kaliappan, A. G. M. S. Raza, P. Anbalagan, and S. Al-Makhadmeh, "Analyzing classification and feature selection strategies for diabetes prediction," *Frontiers in Artificial Intelligence*, vol. 7, p. 1421751, 2024. [Online]. Available: <https://www.frontiersin.org/journals/artificial-intelligence/articles/10.3389/frai.2024.1421751/full>
- [4] M. Kibria, T. A. B. Taz, M. S. Miah, and M. A. Rahman, "An Ensemble Approach for the Prediction of Diabetes Mellitus Using a Soft-Voting Classifier with an Explainable AI," *Sensors*, vol. 22, no. 19, p. 7268, 2022. [Online]. Available: <https://www.mdpi.com/1424-8220/22/19/7268>
- [5] J. Liu, Y. Liu, Y. Xiong, X. Chen, and Y. Liu, "Predicting the Risk of Incident Type 2 Diabetes Mellitus in Chinese Elderly Using Machine Learning Technique," *Frontiers in Public Health*, vol. 10, p. 10107388, 2022. [Online]. Available: <https://pmc.ncbi.nlm.nih.gov/articles/PMC10107388/>
- [6] J. Brownlee, "Feature Importance and Feature Selection With XGBoost in Python," *Machine Learning Mastery*, 2019. [Online]. Available: <https://machinelearningmastery.com/feature-importance-and-feature-selection-with-xgboost-in-python/>
- [7] A. Lundberg and S.-I. Lee, "XGBoost gain-based feature importance vs. SHAP normalized importance (averaged for all patients)," *ResearchGate*, 2021. [Online]. Available: https://www.researchgate.net/figure/XGBoost-gain-based-feature-importance-vs-SHAP-normalized-importance-averaged-for-all_fig3_393302056
- [8] C. Molnar, "Interpretable machine learning: A guide for making black box models explainable," *Computer Methods and Programs in Biomedicine*, vol. 206, p. 106581, 2021. [Online]. Available: <https://www.sciencedirect.com/science/article/abs/pii/S0169260721006581>
- [9] Ghatnekar, Atharva, and Aakash Dhananjay Shanbhag. "Explainable, multi-region price prediction." In 2021 International conference on electrical, computer and energy technologies (ICECET), pp. 1-7. IEEE, 2021.
- [10] Vashi, C., & Shanbhag, A. (2023, July). Comparative explainable Machine learning to evaluate Bank marketing success. In 2023 3rd International Conference on Electrical, Computer, Communications and Mechatronics Engineering (ICECCME) (pp. 1-5). IEEE.