



Science Forum (Journal of Pure and Applied Sciences)

Journal Homepage: <https://atbuscienceforum.com.ng/index.php/jpas>



Explainable Machine Learning for Cardiovascular Disease Detection: Addressing Interpretability Limitations and Evaluating Transparency for Clinical Decision Support

Lateefat Yusuf Saheed^{1,2}, Saratu Yusuf Ilu², Saheed T. Zubair³, Micah Sunom Daniel⁴, Noah N. Gana⁵

¹Department of Software Engineering, Faculty of Science and Computing, Baba Ahmed University Kano, Kano State, Nigeria.

²Department of Computer Science, Faculty of Computing, Bayero University Kano, Kano State, Nigeria.

³Department of Electrical Engineering, Faculty of Engineering, Bayero University Kano, Kano State, Nigeria.

⁴ICT and Data Management Unit, Department of Operations, Kaduna State Residents Identity Management Agency, Kaduna State, Nigeria.

⁵Department of Cyber Security, Nile University of Nigeria, Abuja - Nigeria

Abstract

Cardiovascular diseases (CVDs) remain the leading cause of mortality worldwide, accounting for approximately 17.9 million deaths annually. Although machine learning (ML) models have demonstrated high predictive performance for CVD detection, their inherent black-box nature limits clinical trust, interpretability, and adoption in healthcare practice. This study addresses three objectives: (1) to examine the limitations of conventional ML models for CVD prediction; (2) to develop an interpretable ML framework that enhances model transparency through the integration of SHapley Additive exPlanations (SHAP) and Local Interpretable Model-agnostic Explanations (LIME); and (3) to evaluate the predictive performance, transparency, and clinical applicability of the resulting models. Logistic Regression, Random Forest, and Extreme Gradient Boosting (XGBoost) classifiers were trained and evaluated using the Kaggle Cardiovascular Disease Dataset comprising 70,000 patient records and 11 predictor variables. Model performance was assessed using accuracy, precision, recall, F1-score, receiver operating characteristic area under the curve (ROC-AUC), calibration error, and a structured transparency scoring framework. XGBoost achieved the highest discriminative performance, with an ROC-AUC of 0.794, an accuracy of 0.729, and the lowest calibration error of 0.024, whereas Logistic Regression attained the highest transparency score (73.8/100). Global model interpretation using SHAP identified cholesterol, age, systolic blood pressure, and diastolic blood pressure as the most influential predictors, consistent with established clinical evidence. Similarly, LIME provided patient-specific explanations that highlighted the same dominant risk factors, thereby enhancing decision interpretability at the individual level. These findings demonstrate that high predictive accuracy and model explainability can coexist within a unified ML framework, providing a replicable blueprint for transparent cardiovascular risk assessment, particularly in resource-limited healthcare settings.

Keywords

Cardiovascular disease detection;
Explainable artificial intelligence;
SHAP;
Machine learning transparency;
Clinical decision support



1. INTRODUCTION

Cardiovascular diseases (CVDs), including coronary artery disease, heart failure, and stroke, remain the foremost cause of global morbidity and mortality, claiming approximately 17.9 million lives annually and representing 32% of all deaths worldwide (World Health Organization, 2022). In Nigeria, CVDs account for nearly 30% of all-cause mortality, with hypertension affecting roughly one-third of adults (Okwuosa et al., 2019; Adegbija et al., 2015). This burden is compounded by limited access to specialised diagnostics outside major urban centres, motivating the need for data-driven, low-cost risk stratification tools.

Machine learning has shown strong capacity to detect complex, non-linear patterns in clinical data, with supervised classifiers such as logistic regression, random forests, and gradient boosting frequently outperforming conventional risk scores in CVD prediction tasks, recording benchmark accuracies of 0.70–0.85 across diverse datasets (Shah et al., 2019; Rajkomar et al., 2019; Krittanawong et al., 2020). However, a persistent tension exists between predictive performance and model interpretability: while ensemble models such as random forests and XGBoost typically outperform linear classifiers, they sacrifice the readable decision rules offered by a single decision tree, producing opaque predictions that clinicians cannot interrogate or verify (Rudin, 2019). Additional constraints, including sensitivity to data imbalance and clinically dangerous false-negative rates in screening contexts further limit the safe deployment of conventional ML in healthcare (Chicco and Jurman, 2020).

Explainable Machine Learning (XML) has emerged as a principled response to these limitations. SHAP (Lundberg and Lee, 2017) breaks down individual predictions into additive feature contributions using Shapley values, guaranteeing consistency and local accuracy, while LIME (Ribeiro et al., 2016) constructs locally faithful linear approximations around individual predictions to produce case-specific explanations. Both methods have been applied to cardiovascular risk tasks: SHAP has been used to validate feature hierarchies in heart disease data (Uyar and Ilhan, 2021), and LIME has facilitated individual-level risk communication in CVD classification (Senan et al., 2022).

Broader healthcare applications of XAI have likewise demonstrated its value in improving clinician trust and model adoption (Ogunpola et al., 2024; Srinivasan et al., 2023). However, existing studies typically demonstrate SHAP or LIME outputs descriptively without systematically comparing transparency scores across multiple classifiers, quantifying clinical error distributions, or assessing explanation stability — limitations that constrain the translational utility of prior work.

This study addresses those gaps through a unified, multi-criteria evaluation framework applied to three classifiers on a large-scale CVD dataset. The unique novelty of this work is threefold: first, it introduces a structured transparency scoring framework that quantifies interpretability across four operationalised dimensions (model interpretability, feature-importance availability, decision-rule accessibility, and probability reliability), enabling direct, reproducible comparison of classifier transparency rather than relying on qualitative

description; second, it conducts an Explanation Quality Score (EQS) assessment that formally measures the stability of LIME-generated explanations under repeated sampling, a methodological contribution absent from prior CVD XAI studies; and third, it integrates context-specific clinical error analysis (Type I and Type II rates) with explainability assessment to provide deployment recommendations tailored to distinct clinical use cases. The study therefore pursues three specific objectives: (i) to identify and analyse the interpretability limitations of traditional ML models in CVD detection; (ii) to design an explainable ML framework integrating SHAP and LIME to address these limitations; and (iii) to evaluate the resulting models on transparency, predictive performance, calibration, and clinical-decision-support utility.

2. METHODOLOGY

2.1 Dataset

The publicly available Kaggle Cardiovascular Disease Dataset was used (Ulianova, 2019), comprising 70,000 anonymised patient records with 11 predictor variables and a binary CVD outcome. Features span demographic (age, gender), anthropometric (height, weight), physiological (systolic/diastolic blood pressure, cholesterol, glucose), and behavioural (smoking, alcohol, physical activity) domains. The target is balanced (50% positive/negative). Age was converted from

days to years, and Body Mass Index (BMI) was derived from height and weight.

2.2 Preprocessing and Model Development

A preprocessing pipeline was implemented in Python 3.9 (Pandas, NumPy, Scikit-learn). Steps included descriptive profiling, outlier removal for physiologically implausible blood-pressure readings, BMI categorisation per WHO thresholds, categorical encoding, feature standardisation, and stratified 80:20 train-test partitioning. Three classifiers spanning the interpretability-performance spectrum were developed: Logistic Regression (interpretable baseline), Random Forest (ensemble with feature importance), and XGBoost (state-of-the-art boosting). Hyperparameters were tuned via grid search with stratified 5-fold cross-validation.

2.3 Evaluation and Explainability Framework

Models were evaluated on accuracy, precision, recall, F1-score, ROC-AUC, and calibration error. A structured transparency scoring framework rated each model on interpretability, feature-importance availability, decision-rule accessibility, and probability reliability (each /100). Clinical error analysis quantified Type I and Type II rates from confusion matrices. SHAP provided global (summary, dependence plots) and local (force plot) explanations; LIME generated patient-specific local approximations. An Explanation Quality Score (EQS) assessed LIME consistency across repeated sampling on identical XGBoost predictions.

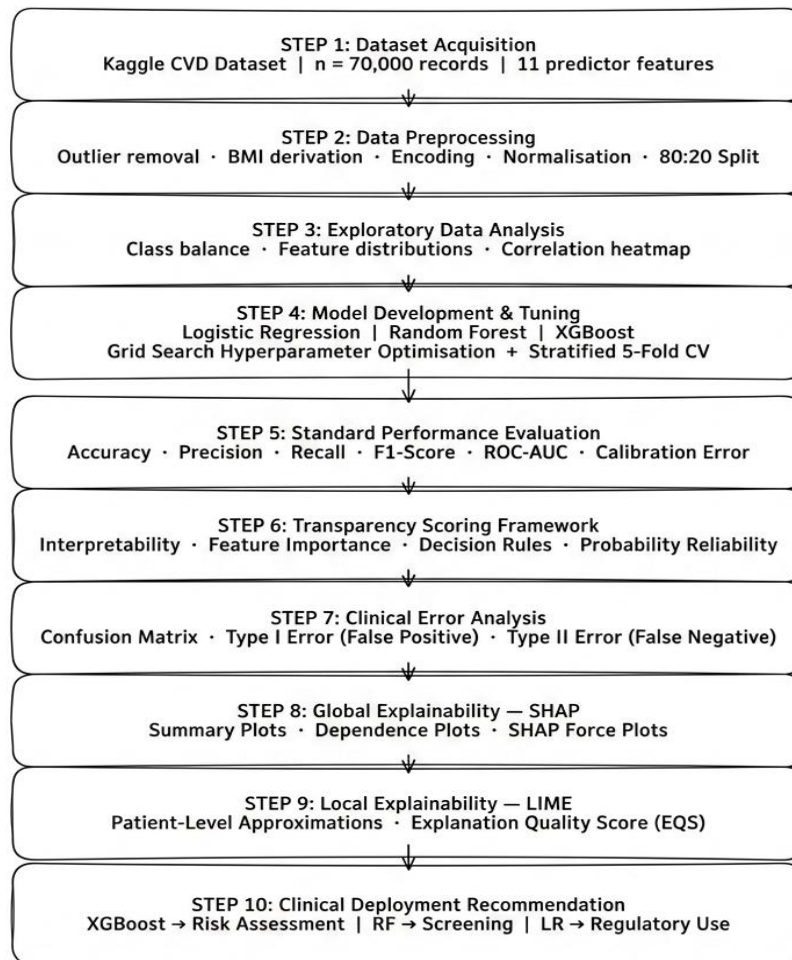


Figure 1: Methodology Flowchart for Explainable Machine Learning for CVD Detection

3. RESULTS AND DISCUSSION

3.1 Limitations of Traditional ML Models

The transparency scoring framework (Table 1) confirmed the interpretability deficit of complex models. XGBoost scored lowest (63.8/100), driven primarily by limited decision-rule accessibility (40/100), reflecting the inherent opacity of sequential boosted tree ensembles. Despite the highest accuracy (0.729), this opacity is clinically problematic: regulatory and ethical frameworks — including the European Union's General Data Protection Regulation (GDPR) Article 22 and the WHO guidance on AI in healthcare — require that automated high-stakes decisions be explainable to affected individuals (Wachter et al., 2017; WHO,

2021). Logistic Regression, while most transparent (73.8/100) owing to its directly interpretable coefficient structure, produced the highest Type II error rate (32.6%), meaning nearly one in three true CVD patients would be missed under a screening deployment — a miss rate that poses unacceptable patient safety risk. This empirically confirms the finding of Rudin (2019) that optimising for a single criterion — accuracy or interpretability alone — is insufficient for safe clinical deployment, and extends that argument by quantifying the trade-off numerically across a large real-world CVD dataset. Random Forest occupied an intermediate position on both transparency (66.2/100) and performance, consistent with prior benchmark comparisons (Ogunpola et al., 2024).

Table 1: Transparency Scoring Framework Across Classifiers

	Transparency (/100)	Interpretability	Feat. Importance	Decision Rules	Probability	Accuracy
Logistic Regression	73.8	90	80	40	85	0.713
Random Forest	66.2	60	80	40	85	0.711
XGBoost	63.8	50	80	40	85	0.729

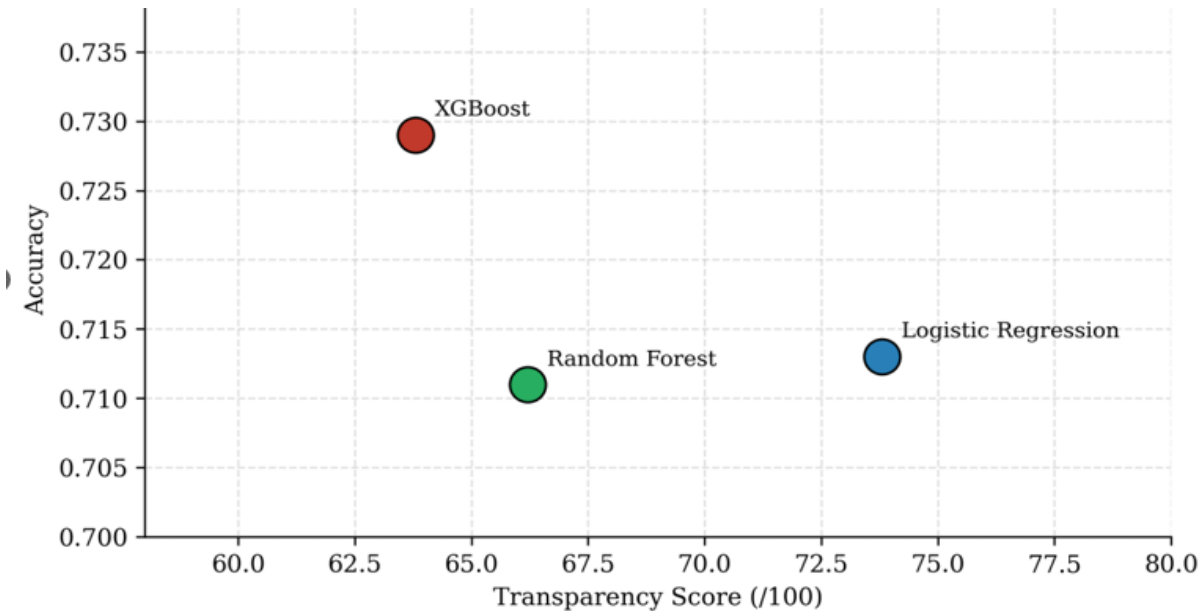


Figure 2: Transparency-Accuracy Trade-off Across Classifiers. XGBoost achieves the highest accuracy but lowest transparency, while Logistic Regression shows the inverse pattern.

3.2Explainable ML Framework Design

Global SHAP analysis for XGBoost (Figure 3) consistently ranked systolic blood pressure as the most influential predictor (18-22% of predictive contribution), followed by age (14-17%), diastolic blood pressure (11-14%), and cholesterol level (9-12%). BMI and glucose contributed moderately (6-9%), while behavioural factors (smoking, alcohol, physical activity) collectively accounted for 10-15%.

This risk factor hierarchy is strongly consistent with established clinical cardiovascular epidemiology: systolic hypertension and advancing age are the two most established modifiable and non-modifiable CVD risk factors respectively (Whelton et al., 2018; Benjamin et al., 2019), while elevated cholesterol is a primary pharmacological intervention target (Powell-Wiley et al., 2021). The alignment of SHAP rankings with clinical knowledge constitutes an

important form of external validation, distinguishing this framework from black-box models whose feature weights cannot be verified clinically (Doshi-Velez and Kim, 2017). Comparable SHAP-based validation has been reported in related cardiovascular studies: Uyar and Ilhan (2021) similarly identified

blood pressure and age as top SHAP contributors in a smaller heart disease dataset, while Baghdadi et al. (2023) confirmed age and physiological markers as dominant predictors in a separate clinical AI application.

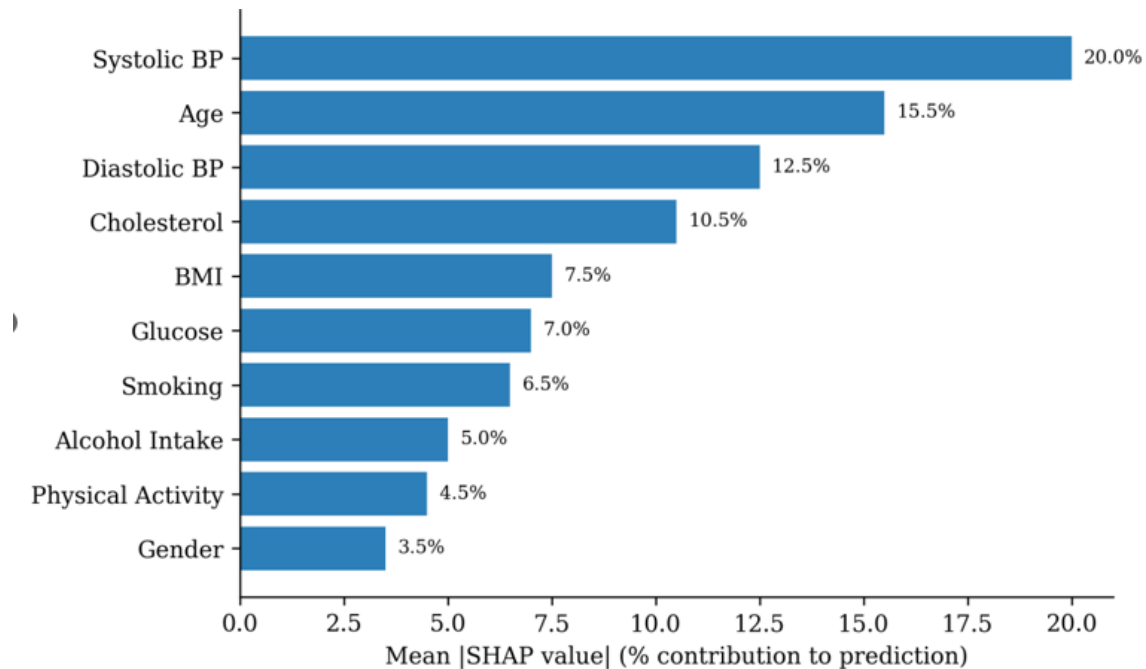


Figure 3: Global SHAP Feature Importance for the XGBoost Model, ranked by mean absolute SHAP value.

SHAP dependence plots further revealed clinically meaningful non-linear threshold effects. Systolic blood pressure showed an accelerated risk increase beyond 140 mmHg — precisely the stage-2 hypertension boundary defined by ACC/AHA guidelines (Whelton et al., 2018) and age exhibited progressive risk escalation after 50 years. These quantified thresholds provide actionable clinical intelligence absent from conventional population-level risk scores such as Framingham or SCORE. At the local level, LIME explanations demonstrated patient-

specific reasoning capacity: two patients sharing an identical 65% predicted CVD probability showed distinct explanation profiles one driven by age and blood pressure, another by cholesterol and BMI illustrating the framework's ability to support truly individualised clinical communication, consistent with the patient-centred explanation goals advocated by Ribeiro et al. (2016).

3.3 Performance, Calibration, and Clinical Utility

Table 2 and Figure 1 summarise comprehensive performance metrics. XGBoost achieved the strongest overall discrimination (ROC-AUC = 0.794; Accuracy = 0.729) and the lowest calibration error (0.024), indicating that its predicted probabilities closely reflect true event rates, a property critical for clinical probability-based risk stratification and patient counselling (Shah et al., 2019). Logistic Regression ranked second on ROC-AUC (0.778) but recorded substantially higher calibration error (0.065), limiting its utility for probability communication despite its transparency advantage. Random Forest achieved the highest recall (0.700),

making it most sensitive to true CVD cases. Relative to prior literature, the XGBoost ROC-AUC of 0.794 is consistent with the 0.78-0.82 range reported by Ogunpola et al. (2024) in a comparable CVD prediction review, and with the 0.79 AUC achieved by Srinivasan et al. (2023) on a UCI heart disease dataset, validating the competitiveness of the proposed framework. Logistic Regression's AUC of 0.778 aligns with the 0.76-0.80 range documented for linear models in CVD risk prediction (Rajkomar et al., 2019; Krittanawong et al., 2020), while Random Forest's AUC of 0.769 is consistent with Baghdadi et al. (2023) who reported 0.75-0.80 for ensemble methods on physiological datasets.

Table 2: Comprehensive Model Performance Comparison

Model	Accuracy	Precision	Recall	F1	ROC-AUC	Calib. Error
XGBoost	0.729	0.746	0.692	0.718	0.794	0.024
Logistic Regression	0.713	0.731	0.674	0.701	0.778	0.065
Random Forest	0.711	0.716	0.700	0.708	0.769	0.048

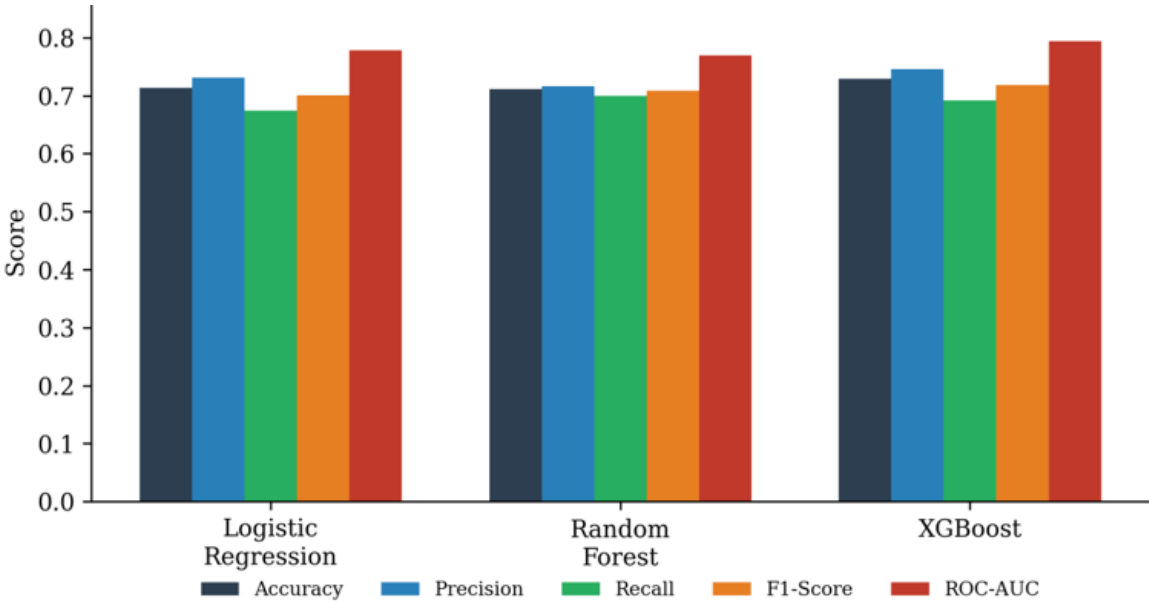


Figure 4: Model Performance Comparison Across Accuracy, Precision, Recall, F1-Score, and ROC-AUC.

Clinical error analysis (Table 3) further differentiates model utility by deployment context. XGBoost produced the lowest Type I error rate (23.5%), minimising false alarms and the downstream burden of unnecessary investigations — a practical consideration in resource-constrained Nigerian primary healthcare facilities where over-referral imposes substantial costs. Random Forest achieved the lowest Type II error rate (30.0%), maximising true CVD case detection and reducing the patient safety risk of missed

diagnoses. Logistic Regression recorded the highest Type II error rate (32.6%), rendering it unsuitable as a standalone screening instrument despite its transparency advantage. These differential error profiles mirror findings by Chicco and Jurman (2020), who demonstrated that optimising for overall accuracy conceals clinically relevant Type II error variation across classifiers in cardiovascular datasets, underscoring the importance of clinical error profiling alongside standard metrics.

Table 3: Clinical Error Analysis (Test Set $n \approx 14,000$)

Metric	Logistic Regression	Random Forest	XGBoost
True Positives	4,715	4,900	4,844
True Negatives	5,270	5,058	5,355
False Positives	1,734	1,946	1,649
False Negatives	2,281	2,096	2,152
Type I Error Rate	24.8%	27.8%	23.5%
Type II Error Rate	32.6%	30.0%	30.8%

The Explanation Quality Score for XGBoost under LIME was notably low (EQS = 0.26): repeated sampling on the same patient record produced divergent explanations, attributing high risk to different features across runs. This instability, consistent with the approximation-variance limitations of LIME documented by Slack et al. (2020) and Alvarez-Melis and Jaakkola (2018), indicates that LIME outputs for complex, non-linear models such as XGBoost should be presented to clinicians with explicit uncertainty qualification rather than as definitive explanations. The deterministic SHAP force plot which produced consistent, mathematically grounded attributions is recommended as the primary explanation

interface for clinical deployment. Logistic Regression's coefficients, serving directly as globally consistent explanations, achieved the highest EQS by construction, reinforcing its suitability where regulatory-grade interpretability is mandated. Overall, the results demonstrate that the accuracy-transparency trade-off is not a binary constraint but a spectrum navigable through a layered framework: XGBoost for risk stratification with SHAP-based explanation, Random Forest for high-sensitivity screening, and Logistic Regression for regulatory-compliant transparent deployment.

4. CONCLUSION

This study designed and evaluated an explainable machine learning framework for cardiovascular disease detection, addressing three objectives: identifying interpretability limitations of traditional ML models, designing a dual-layer SHAP/LIME explainability framework, and evaluating the resulting models' transparency, performance, calibration, and clinical utility. For a balanced dataset (70k records), XGBoost had the best discrimination (ROC-AUC = 0.794) and calibration (error = 0.024) while Logistic Regression had the highest inherent transparency (73.8/100).

SHAP analysis gave reliable, clinically substantiated interpretations, with systolic blood pressure, age and cholesterol as the key risk factors, and LIME-based testing found that XGBoost's stability boundeared its usefulness for clinical deployment, suggesting to use SHAP for that purpose. The framework shows that high predictive ability and valuable clinical explanations could be achieved in a single pipeline, providing a template for the evaluation of future medical AI systems.

Limitations include cross-sectional retrospective data, the lack of clinical data from Nigeria and the exclusion of deep learning architectures. To achieve prospective validation in Nigeria's primary healthcare facilities, as well as longitudinal modelling of electronic health records and fairness-aware assessment of bias across demographic subgroups, future research must be continued.

REFERENCES

- Adegbiya, O., Hoy, W. and Wang, Z. (2015). Prediction of cardiovascular disease risk using waist circumference among Aborigines in a remote Australian community. *BMC Public Health*, 15, 57. <https://doi.org/10.1186/s12889-015-1406-1>
- Alvarez-Melis, D. and Jaakkola, T. S. (2018). On the robustness of interpretability methods. *ICML Workshop on Human Interpretability in Machine Learning*, pp. 1-6. <https://doi.org/48550/arXiv.1806.08049>
- Baghdadi, N. A., Malki, A., Balaha, H. M., AbdulAzeem, Y., Badawy, M. and Elhosseini, M. (2023). An optimized deep learning approach for suicide detection through Arabic tweets. *PeerJ Computer Science*, 9, e1070. <https://doi.org/10.7717/peerj-cs.1070>
- Benjamin, E. J., Muntner, P., Alonso, A., Bittencourt, M. S., Callaway, C. W. and Carson, A. P. (2019). Heart disease and stroke statistics — 2019 update: a report from the American Heart Association. *Circulation*, 139(10), e56-e528. <https://doi.org/10.1161/CIR.00000000000000659>
- Chicco, D. and Jurman, G. (2020). Machine learning can predict survival of patients with heart failure from serum creatinine and ejection fraction alone. *BMC Medical Informatics and Decision Making*, 20(1), 16. <https://doi.org/10.1186/s12911-020-1023-5>

- Doshi-Velez, F. and Kim, B. (2017). Towards a rigorous science of interpretable machine learning. *arXiv preprint*. <https://doi.org/48550/arXiv.1702.08608>
- Krittananawong, C., Johnson, K. W., Rosenson, R. S., Wang, Z., Aydar, M., Baber, U., Min, J. K., Tang, W. W., Halperin, J. L. and Bhatt, D. L. (2020). Deep learning for cardiovascular medicine: a practical primer. *European Heart Journal*, 40(25), 2058-2073. <https://doi.org/10.1093/eurheartj/ehz056>
- Lundberg, S. M. and Lee, S. I. (2017). A unified approach to interpreting model predictions. *Advances in Neural Information Processing Systems*, 30, 4765-4774. <https://proceedings.neurips.cc/paper/2017/hash/8a20a8621978632d76c43dfd28b67767-Abstract.html>
- Ogunpola, A., Saeed, F., Basurra, S., Albarrak, A. M. and Qasem, S. N. (2024). Machine learning-based predictive models for detection of cardiovascular diseases. *Diagnostics*, 14(2), 144. <https://doi.org/10.3390/diagnostics14020144>
- Okwuosa, I. S., Lewsey, S. C., Adesiyun, T., Hicks, A. J. and Mazimba, S. (2019). Cardiovascular disease in sub-Saharan Africans in the diaspora: lessons learned and opportunities to improve care. *JACC: Case Reports*, 1(4), 536-539. <https://doi.org/10.1016/j.jaccas.2019.08.002>
- Powell-Wiley, T. M., Poirier, P., Burke, L. E., Despres, J. P., Gordon-Larsen, P., Lavie, C. J., Lear, S. A., Ndumele, C. E., Neeland, I. J. and Sanders, P. (2021). Obesity and cardiovascular disease: a scientific statement from the American Heart Association. *Circulation*, 143(21), e984-e1010. <https://doi.org/10.1161/CIR.00000000000000973>
- Rajkomar, A., Dean, J. and Kohane, I. (2019). Machine learning in medicine. *New England Journal of Medicine*, 380(14), 1347-1358. <https://doi.org/10.1056/NEJMr1814259>
- Ribeiro, M. T., Singh, S. and Guestrin, C. (2016). 'Why should I trust you?': Explaining the predictions of any classifier. *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, pp. 1135-1144. <https://doi.org/10.1145/2939672.2939778>
- Rudin, C. (2019). Stop explaining black box machine learning models for high stakes decisions and use interpretable models instead. *Nature Machine Intelligence*, 1(5), 206-215. <https://doi.org/10.1038/s42256-019-0048-x>
- Senan, E. M., Barhoom, A. H., Al-Mekhlafi, Z. G., Alreshidi, A., Abd El-Latif, A. A., Ibrahim, D. M. and Aly, M. H. (2022). Diagnosis of chronic kidney disease using effective classification algorithms and recursive feature elimination techniques. *Journal of Healthcare Engineering*, 2022, Article ID 4450389. <https://doi.org/10.1155/2022/4450389>

- Shah, S. J., Katz, D. H., Selvaraj, S., Burke, M. A., Yancy, C. W., Gheorghiade, M., Bonow, R. O., Huang, C. C. and Deo, R. C. (2015). Phenomapping for novel classification of heart failure with preserved ejection fraction. *Circulation*, 131(3), 269-279. <https://doi.org/10.1161/CIRCULATIONAHA.114.010637>
- Slack, D., Hilgard, S., Jia, E., Singh, S. and Lakkaraju, H. (2020). Fooling LIME and SHAP: adversarial attacks on post hoc explanation methods. *Proceedings of the AAAI/ACM Conference on AI, Ethics, and Society*, pp. 180-186. <https://doi.org/10.1145/3375627.3375830>
- Srinivasan, S., Gunasekaran, S., Mathivanan, S. K., Muthukumaran, V., Babu, J. C. and Herencsar, N. (2023). An active learning machine technique-based prediction of cardiovascular heart disease from UCI-repository. *Scientific Reports*, 13(1), 13588. <https://doi.org/10.1038/s41598-023-40717-1>
- Ulianova, S. (2019). Cardiovascular disease dataset. *Kaggle*. <https://www.kaggle.com/datasets/sulianova/cardiovascular-disease-dataset>
- Uyar, K. and Ilhan, A. (2021). Diagnosis of heart disease using genetic algorithm-based trained recurrent fuzzy neural networks. *Procedia Computer Science*, 120, 588-594. <https://doi.org/10.1016/j.procs.2017.11.283>
- Wachter, S., Mittelstadt, B. and Russell, C. (2017). Counterfactual explanations without opening the black box: automated decisions and the GDPR. *Harvard Journal of Law and Technology*, 31(2), 841-887. <https://doi.org/48550/arXiv.1711.00399>
- Whelton, P. K., Carey, R. M., Aronow, W. S., Casey, D. E., Collins, K. J., Dennison Himmelfarb, C. et al. (2018). 2017 ACC/AHA guideline for the prevention, detection, evaluation, and management of high blood pressure in adults. *Journal of the American College of Cardiology*, 71(19), e127-e248. <https://doi.org/10.1016/j.jacc.2017.11.006>
- World Health Organization (2021). Ethics and governance of artificial intelligence for health. *WHO Guidance*. World Health Organization, Geneva. <https://www.who.int/publications/i/item/9789240029200>
- World Health Organization (2022). Cardiovascular diseases (CVDs). *WHO Fact Sheet*. World Health Organization, Geneva. [https://www.who.int/news-room/fact-sheets/detail/cardiovascular-diseases-\(cvds\)](https://www.who.int/news-room/fact-sheets/detail/cardiovascular-diseases-(cvds))