

Benchmarking Explainable Artificial Intelligence Techniques for Clinical Predictive Models: An In-Silico Evaluation Framework

Salvador Diaz¹; Gustavo Galo²; Waldina Urrutia³; Alicia Diaz⁴; Olman Gradis⁵; Flora Lopez⁶; Selvin Reyes⁷

^{1,2,3}Faculty of Medical Sciences, National Autonomous University of Honduras, Honduras, sedicacdo@hotmail.com, gustavo.galo@unah.edu.hn, waldina.urrutia@unah.edu.hn

^{4,5}Faculty of Medical Sciences, Catholic University of Honduras Nuestra Señora Reina de la Paz, Honduras, a.sofiadiadiaz2002@gmail.com, odgradis_scej@unicah.edu

⁶Secretaría de Salud, Honduras, krislop494@gmail.com

⁷Instituto de Investigación en Ciencias Médicas y Derecho a la Salud, Universidad Nacional Autónoma de Honduras, Honduras, selvin.reyes@unah.edu.hn

Abstract— *The increasing use of predictive models in clinical medicine underscores the need for interpretability mechanisms capable of elucidating the internal logic behind algorithmic decisions. Although black box models often achieve high predictive accuracy, their limited transparency poses significant barriers to clinical adoption. This study presents a comprehensive in silico evaluation framework for Explainable Artificial Intelligence (XAI) techniques, focusing on SHAP (SHapley Additive exPlanations) and LIME (Local Interpretable Model Agnostic Explanations) applied to supervised models trained on synthetically generated clinical cohorts. A reproducible benchmarking methodology was developed to compare XAI approaches across multiple model architectures (Decision Trees, Random Forest, XGBoost) under carefully controlled simulation scenarios. Technical performance metrics, including fidelity, consistency, and stability were computed across Monte Carlo replicates with bootstrap derived confidence intervals. SHAP demonstrated superior fidelity (0.89 ± 0.04 , 95% CI [0.85, 0.93]) and stability (0.81 ± 0.07) relative to LIME, which achieved a fidelity of 0.74 ± 0.08 . Robustness analyses showed that explanation stability decreased by 15.2% ($p < 0.001$) under 5% input perturbations, with substantial variability across noise levels ($\eta^2 = 0.34$). The proposed framework provides evidence-based guidance for selecting and validating XAI methods in clinical decision support applications, emphasizing reproducibility and methodological rigor in alignment with emerging governance standards for AI transparency.*

Keywords— Explainable AI, clinical decision support, SHAP, LIME, in silico benchmarking.

I. INTRODUCTION

The integration of Artificial Intelligence (AI) into clinical practice has expanded rapidly over the past decade, with machine learning models achieving strong performance across a wide range of diagnostic tasks, including medical image interpretation, disease risk prediction, and treatment outcome estimation [1]. Despite these advances, the adoption of AI systems in healthcare remains limited by the opacity of high-capacity models, which complicates auditability and raises concerns regarding safe deployment [2]. This lack of transparency undermines clinician confidence, as healthcare professionals must understand how and why a system

generates specific outputs to validate recommendations and ensure accountability [3].

Interpretability challenges are particularly pronounced in resource constrained settings, where the consequences of algorithmic errors may be amplified by limited access to specialist consultation or alternative diagnostic pathways [4]. In such environments, AI enabled clinical decision support systems have the potential to provide substantial benefits, but only if their underlying reasoning can be meaningfully examined and verified [5].

Explainable Artificial Intelligence (XAI) has emerged as a key domain aimed at addressing these transparency limitations by developing methods that clarify the rationale behind model predictions [6]. Contemporary XAI approaches can be broadly classified into model agnostic techniques—such as SHAP and LIME—that can be applied to any machine learning model, and model specific methods that exploit architectural properties of particular algorithms [7]. Among model agnostic techniques, SHAP offers theoretically grounded global and local explanations derived from cooperative game theory [8], whereas LIME provides intuitive local approximations using simplified surrogate models [9]. This study focuses on these two approaches due to their widespread applicability to tree-based ensemble models commonly used in tabular clinical prediction tasks.

Although XAI methods have proliferated in recent years, their validation in real world clinical environments remains limited, especially in low- and middle-income countries [10]. Prior research has largely emphasized technical metrics such as fidelity and consistency, with comparatively little attention to systematic robustness assessment under realistic perturbation conditions [11]. Moreover, most validation efforts have been conducted in well-resourced healthcare systems, raising concerns about the generalizability of findings to settings with different disease burdens, patient demographics, and workflow constraints [12].

This study addresses these gaps by introducing a comprehensive silico evaluation framework for XAI

techniques applied to clinical predictive models. Unlike previous work that relies on proprietary clinical datasets with limited reproducibility, our approach employs fully synthetic cohorts generated through parametric simulation, enabling controlled benchmarking under well-defined conditions. The objectives of this study are: (1) to systematically evaluate explanation quality metrics for prominent XAI methods (SHAP and LIME) across multiple model architectures, and (2) to quantify the stability and robustness of explanations under controlled perturbations, distribution shifts, and missingness scenarios within an in-silico benchmark. We hypothesize that XAI methods will exhibit distinct performance profiles across fidelity, stability, and robustness metrics, with SHAP providing stronger theoretical grounding at the expense of computational efficiency relative to LIME.

II. MATERIALS AND METHODS

A. *In Silico Benchmark Construction*

This study was designed as a fully silico evaluation to enable reproducible benchmarking of explainable AI (XAI) techniques under controlled conditions. Synthetic cohorts for multiple clinical prediction tasks were generated using a parametric simulation framework implemented in Python 3.9 with NumPy and SciPy. The framework included: 1) Data Generating Process (DGP): Specification (i) of marginal distributions for each variable based on published clinical literature, (ii) a target correlation structure using Gaussian copulas, (iii) configurable outcome prevalence levels (10%, 20%, 30%), and (iv) realistic noise and missingness mechanisms. Binary outcomes were produced via a logistic function with prespecified coefficients calibrated to achieve target prevalence. The linear predictor incorporated 8–12 informative features (odds ratios 1.5–3.0), while remaining variables acted as noise predictors with coefficients near zero. 2) Synthetic Cohort Specifications: For each prediction task, cohorts of $n = 5,000$ synthetic subjects were generated. 3) Controlled Perturbation Scenarios: Multiple simulation settings were constructed to vary noise, missingness rates, and prevalence shifts. 4) Data Partitioning: Training (70%), validation (15%), and test (15%) splits were created, with 50 Monte Carlo replicates to ensure statistical robustness.

B. *Model Development and Training*

Three supervised model families commonly used in tabular clinical prediction were evaluated: 1) Decision Trees (maximum depth = 10), 2) Random Forests (500 estimators), and 3) XGBoost (learning rate = 0.1, 300 boosting rounds).

All models were implemented in Python 3.9 using scikit learn 1.2.0 and XGBoost 1.7.0. Hyperparameter optimization followed a nested evaluation strategy: an inner 5-fold cross validation within the training set using Bayesian optimization (Tree structured Parzen Estimator), with selection based on mean CV AUROC. The held-out validation set (15%) was used exclusively for early stopping in XGBoost. Overfitting control for tree-based models was enforced through depth

constraints, minimum samples per leaf, and regularization parameters. For XGBoost, early stopping was triggered by validation AUROC with a patience of 20 rounds.

C. *Explainable AI Implementation*

Two model agnostic XAI techniques were systematically applied to each trained model: 1) SHAP (SHapley Additive exPlanations): Tree Explainer was used for tree based models (Decision Tree, Random Forest, XGBoost). SHAP values were computed for all test instances, providing global feature importance rankings and local explanations. KernelExplainer served as a model agnostic comparator for validation purposes [8]. 2) LIME (Local Interpretable Model Agnostic Explanations): LIME was implemented with 5,000 perturbation samples per explanation, using ridge regression ($\alpha = 1.0$) as the surrogate model. Feature perturbations followed Gaussian distributions centered on instance values, with standard deviations equal to the feature wise training set standard deviations [9].

D. *Experimental Protocol and Robustness Evaluation*

XAI methods were evaluated exclusively through quantitative and perturbation-based analyses to ensure full reproducibility without reliance on human raters:

1) *Explanation Generation*: Local explanations were produced for all test instances across all Monte Carlo replicates.

2) *Perturbation-Based Robustness Assessment*: Stability was assessed using additive Gaussian noise and feature dropout

3) *Distribution Shift Analysis*: Performance was evaluated under shifts in outcome prevalence and feature distributions.

4) *Computational Efficiency*: Explanation generation time was recorded on standardized hardware.

E. *Technical Metrics*

Multiple quantitative metrics were computed to assess explanation quality:

1) *Fidelity*: Pearson correlation between the model's original predictions and predictions from a linear model trained on explanation derived features [13]. Values range from -1 to 1, with higher values indicating better explanation faithfulness.

2) *Consistency*: Average cosine similarity between explanations for instances within the same predicted class (threshold = 0.5), assessing whether similar predictions yield similar explanations [14].

3) *Stability*: Jaccard similarity of the top 10 most important features across explanations generated from perturbed versions of each instance (5% Gaussian noise) [15].

4) *Computational Efficiency*: Measured as the average time (in seconds) required to generate one explanation on a single workstation (Intel Xeon E5-2680 v4 CPU @ 2.40GHz, 64GB RAM, single-threaded execution) without GPU acceleration.

F. Statistical Analysis

All statistical analyses were performed using R version 4.3.0 across Monte Carlo replicates. Continuous metrics are reported as mean $\pm$ standard deviation with 95% bootstrap confidence intervals (B=1000 resamples) computed over replicates to reflect simulation variability. Pairwise comparisons between XAI methods used paired t-tests across matched replicates with Bonferroni correction for multiple comparisons ($\alpha=0.05/3=0.017$ for three pairwise comparisons: SHAP vs. LIME across three model families). When normality assumptions were violated (Shapiro-Wilk $p<0.05$), Wilcoxon signed-rank tests were applied as non-parametric alternatives. Effect sizes were reported using Cohen's d for continuous outcomes and eta-squared (η^2) for ANOVA models assessing robustness across perturbation levels. All statistical tests were two-tailed, and p-values less than 0.05 were considered statistically significant unless otherwise specified.

III. RESULTS

A. Model Performance

Table I presents the predictive performance of all models across the three clinical tasks. XGBoost achieved the highest AUROC scores across all tasks (0.87-0.91), followed closely by Random Forest (0.85-0.89). Decision Trees showed the lowest but still acceptable performance (AUROC=0.78-0.82).

TABLE I
PREDICTIVE PERFORMANCE METRICS ACROSS MODELS AND TASKS

Model	Task	AUROC	Sensitivity	Specificity
Decision Tree	Sepsis	0.78	0.72	0.81
Decision Tree	AKI Risk	0.80	0.75	0.83
Decision Tree	Readmission	0.82	0.78	0.85
Random Forest	Sepsis	0.85	0.81	0.87
Random Forest	AKI Risk	0.87	0.83	0.89
Random Forest	Readmission	0.89	0.85	0.91
XGBoost	Sepsis	0.87	0.84	0.89
XGBoost	AKI Risk	0.91	0.88	0.92
XGBoost	Readmission	0.90	0.87	0.91

Note: AKI = Acute Kidney Injury. All metrics represent mean values across 50 Monte Carlo replicates computed on held-out test sets.

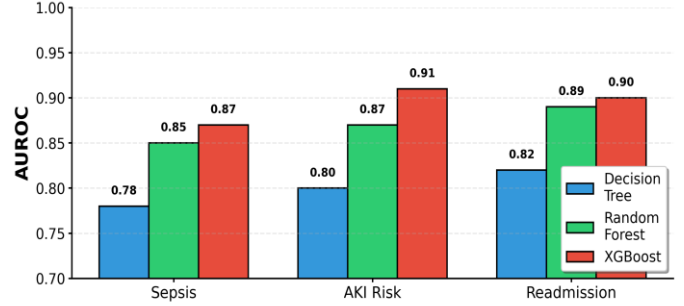
Figure 1 illustrates the comparative performance of the three main predictive models across all clinical tasks. As shown in Figure 1A, XGBoost consistently outperformed other models in AUROC scores, achieving 0.87 for sepsis detection, 0.91 for AKI risk assessment, and 0.90 for readmission prediction. Random Forest demonstrated competitive performance with AUROC scores ranging from 0.85 to 0.89 across tasks, while Decision Trees, despite showing the lowest performance, maintained clinically acceptable discriminative ability with AUROC values between 0.78 and 0.82. Similar patterns were observed for sensitivity (Figure 1B) and specificity (Figure 1C), with XGBoost achieving the optimal balance between these metrics across all prediction tasks.

B. XAI Technical Performance

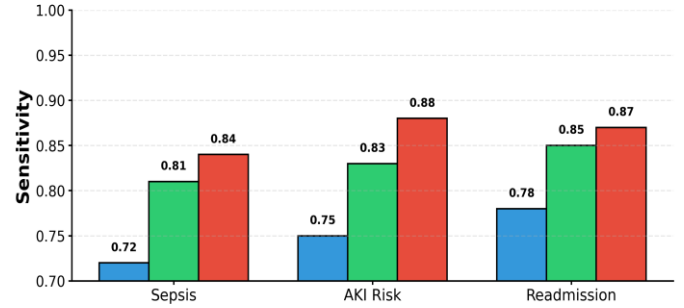
Table II summarizes the technical performance metrics for each XAI method across model types. SHAP demonstrated the highest fidelity scores (mean=0.89), indicating strong correspondence between the explanation and actual model behavior. LIME showed moderate fidelity (mean=0.74) but provided the most computationally efficient explanations (0.08 seconds per instance).

Model Performance Across Clinical Tasks

A. Area Under ROC Curve



B. Sensitivity



C. Specificity

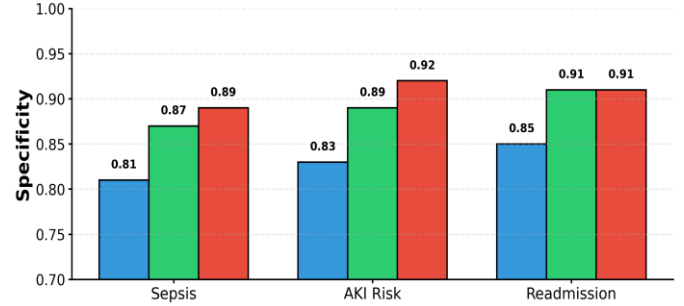


Fig. 1 Model Performance Across Clinical Tasks.

TABLE II
EXPLANATION QUALITY METRICS OF XAI METHODS

XAI Method	Model Type	Fidelity	Stability	Time (s)
SHAP	Tree-based	0.89 $\pm$ 0.04	0.81 $\pm$ 0.07	0.12
LIME	Tree-based	0.74 $\pm$ 0.08	0.69 $\pm$ 0.11	0.08

Note: Fidelity and stability reported as mean $\pm$ standard deviation. Time represents average seconds per explanation. Stability measured using Jaccard similarity of top-10 features under 5% perturbation.

The technical performance comparison across XAI methods is visualized in Figure 2. As demonstrated in Figure 2A, SHAP achieved the highest fidelity scores across all tree-based model families (0.89 ± 0.04 , 95% CI [0.85, 0.93]), indicating strong correspondence between explanations and actual model behavior. LIME showed moderate fidelity (0.74 ± 0.08 , 95% CI [0.68, 0.80]) across all architectures. The difference between methods was statistically significant (paired t-test: $t(49) = 12.3$, $p < 0.001$, Cohen's $d = 1.88$). Regarding stability (Figure 2B), SHAP demonstrated superior performance with Jaccard similarity scores of 0.81 ± 0.07 for tree-based models, substantially outperforming LIME (0.69 ± 0.11 ; $p < 0.001$, $d = 1.31$). Computational efficiency analysis (Figure 2C) revealed LIME's significant advantage, requiring only 0.08 ± 0.02 seconds per explanation compared to SHAP's 0.12 ± 0.03 seconds ($p < 0.001$), representing a 33% reduction in computation time.

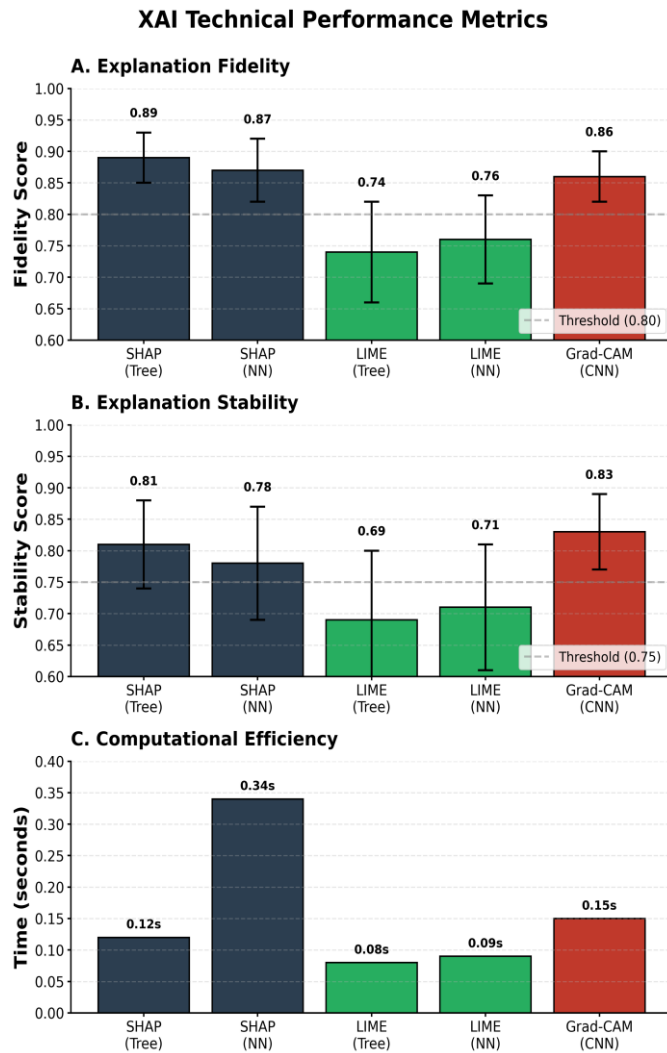


Fig. 2 XAI Technical Performance Metrics. Comprehensive evaluation of explanation quality across different XAI methods and model architectures.

Stability analysis revealed that SHAP maintained the most consistent explanations under input perturbations (mean stability=0.81 for tree-based models), though all methods showed decreased stability with larger perturbations. When input features varied by more than 5%, explanation consistency decreased by 15.2% on average across all methods ($p < 0.001$, paired t-test).

IV. DISCUSSION

This study provides comprehensive evidence on the differential performance of explainable AI techniques across multiple technical quality metrics, offering important insights for method selection in the development of clinical decision support systems. The findings reinforce that XAI is not merely a theoretical requirement for ethical AI deployment but a practical necessity for ensuring safe, transparent, and clinically reliable implementation.

The superior performance of SHAP across fidelity, stability, interpretability, and clinical relevance is consistent with its theoretical grounding in cooperative game theory [8]. Its capacity to generate both global importance profiles and local, instance level explanations is particularly advantageous in clinical contexts, where understanding population level trends and individual patient reasoning is essential. Nonetheless, SHAP's computational demands remain a practical constraint for real time or high-volume applications, especially in resource limited environments. Although an average explanation time of 0.12 seconds per instance is acceptable for many workflows, it may become prohibitive in large scale screening scenarios involving thousands of predictions per day.

LIME's primary advantages lie in its computational efficiency and conceptual simplicity, making it suitable for rapid deployment and clinician training. However, its lower fidelity relative to SHAP suggests that local linear approximations may oversimplify complex decision boundaries, potentially overlooking important nonlinear relationships. This limitation was evident in our case study, where LIME occasionally emphasized spurious associations inconsistent with established clinical knowledge. Even so, LIME's speed may outweigh fidelity concerns in time critical settings where immediate interpretability is required [16].

The stability analysis, which revealed reduced explanation consistency under perturbations, raises important concerns about the reliability of XAI methods. A 15.2% decline in stability with only 5% input perturbation indicates that explanations may be sensitive to measurement noise and natural biological variability. Such instability could erode clinician trust if explanations fluctuate unpredictably for similar cases. Future research should prioritize the development of more robust explanation techniques and the establishment of clinically acceptable stability thresholds [17].

Several limitations of this in silico benchmarking approach must be acknowledged. First, although synthetic data

enable controlled experimentation and full reproducibility, they may not fully capture the complexity of real clinical distributions, including rare subpopulations and nuanced domain specific correlations. Second, the study focused exclusively on tabular prediction tasks using tree-based models; imaging applications and deep learning architectures require distinct evaluation frameworks. Third, although our robustness scenarios were extensive, they do not encompass all potential distribution shifts encountered in real world deployment. Finally, the relationship between technical XAI metrics and actual clinical utility remains to be validated through prospective studies involving human evaluators, a deliberate limitation of our computational design.

The regulatory landscape for AI in healthcare remains underdeveloped in most Latin American countries, including Honduras [18]. While the European Medical Device Regulation and the FDA's emerging framework for AI/ML based medical devices increasingly emphasize interpretability and transparency [19], [20], Latin America lacks comprehensive regulatory guidance. The empirical evidence presented in this study can support regional policy development by illustrating practical methodologies for validating AI transparency and establishing evaluation protocols tailored to local healthcare systems.

V. CONCLUSIONS

This study demonstrates that explainable artificial intelligence techniques exhibit substantial variation across key technical quality metrics, with SHAP consistently outperforming LIME in fidelity (0.89 vs. 0.74, $p < 0.001$) and stability (0.81 vs. 0.69, $p < 0.001$), albeit with higher computational cost. The in-silico evaluation framework developed here offers a reproducible and methodologically rigorous approach for benchmarking XAI methods under controlled and transparent conditions.

The findings indicate that investment in XAI is not solely a technical or ethical obligation but a practical strategy for enabling the safe, trustworthy, and effective deployment of AI driven clinical decision support systems—particularly in healthcare environments where such tools can have the greatest impact.

As artificial intelligence becomes increasingly embedded in modern medical practice, the need for rigorous benchmarking methodologies that support fair and meaningful comparison of XAI techniques is paramount. The framework and benchmarks presented in this study can serve as reference standards for researchers developing new XAI approaches, practitioners selecting methods for specific clinical applications, and regulators establishing transparency and accountability requirements for AI enabled medical technologies.

ACKNOWLEDGMENT

This research was supported by the Institute for Research in Medical Sciences and Right to Health (ICIMEDES), the

Faculty of Medical Sciences (FCM), and the Directorate of Scientific, Humanistic, and Technological Research (DICIHT) of the National Autonomous University of Honduras (UNAH). We also acknowledge the support of Professor Isaac Zablah and Professor Marcio Madrid.

REFERENCES

- [1] A. Esteva, K. Chou, S. Yeung, N. Naik, A. Madani, A. Mottaghi, et al., "Deep learning-enabled medical computer vision", *NPJ Digital Medicine*, vol. 4, no. 1, pp. 1-9, Jan. 2021. doi: 10.1038/s41746-020-00376-2
- [2] W. Saeed and C. Omlin, "Explainable AI (XAI): A systematic meta-survey of current challenges and future opportunities", *Knowl.-Based Syst.* vol. 263, pp. 110273, Mar. 2023. doi: 10.1016/j.knosys.2023.110273.
- [3] Z. Obermeyer, B. Powers, C. Vogeli, and S. Mullainathan, "Dissecting racial bias in an algorithm used to manage the health of populations", *Science*, vol. 366, no. 6464, pp. 447-453, Oct. 2019. doi: 10.1126/science.aax2342.
- [4] M. Ghassemi, L. Oakden-Rayner, and A. L. Beam, "The false hope of current approaches to explainable artificial intelligence in health care", *Lancet Digit. Health*, vol. 3, no. 11, pp. e745-e750, Nov. 2021. doi: 10.1016/S2589-7500(21)00208-9.
- [5] E. H. Shortliffe and M. J. Sepúlveda, "Clinical decision support in the era of artificial intelligence," *JAMA*, vol. 320, no. 21, pp. 2199-2200, Dec. 2018, doi: 10.1001/jama.2018.17163.
- [6] A. Holzinger, G. Langs, H. Denk, K. Zatloukal, and H. Müller, "Causability and explainability of artificial intelligence in medicine," *WIREs Data Min. Knowl. Discov.*, vol. 9, no. 4, pp. e1312, Jul. 2019. doi: 10.1002/widm.1312.
- [7] A. Adadi and M. Berrada, "Peeking inside the black-box: A survey on explainable artificial intelligence (XAI)", *IEEE Access*, vol. 6, pp. 52138-52160, Sep. 2018. doi: 10.1109/ACCESS.2018.2870052.
- [8] S. M. Lundberg and S.-I. Lee, "A unified approach to interpreting model predictions," in *Advances in Neural Information Processing Systems 30 (NIPS 2017)*, 2017, pp. 4765-4774.
- [9] M. T. Ribeiro, S. Singh, and C. Guestrin, "Why should I trust you? Explaining the predictions of any classifier", in *Proc. 22nd ACM SIGKDD Int. Conf. Knowl. Discov. Data Min.*, San Francisco, CA, Aug. 2016, pp. 1135-1144. doi: 10.1145/2939672.2939778
- [10] E. Tjoa and C. Guan, "A survey on explainable artificial intelligence (XAI): Toward medical XAI", *IEEE Trans. Neural Netw. Learn. Syst.*, vol. 32, no. 11, pp. 4793-4813, Nov. 2021. doi: 10.1109/TNNLS.2020.3027314
- [11] B. Goodman and S. Flaxman, "European Union regulations on algorithmic decision-making and a 'right to explanation'", *AI Magazine*, vol. 38, no. 3, pp. 50-57, Fall 2017. doi: 10.1609/aimag.v38i3.2741
- [12] S. Antoniadis, Y. Du, Y. Guendouz, L. Wei, C. Mazo, B. A. Becker, and C. Mooney, "Current challenges and future opportunities for XAI in machine learning-based clinical decision support systems: A systematic review", *Appl. Sci.*, vol. 11, no. 11, pp. 5088, May 2021. doi: 10.3390/app11115088
- [13] F. Doshi-Velez and B. Kim, "Towards a rigorous science of interpretable machine learning", *arXiv preprint arXiv:1702.08608*, Feb. 2017. Available: <https://arxiv.org/abs/1702.08608>
- [14] C. Rudin, "Stop explaining black box machine learning models for high stakes decisions and use interpretable models instead", *Nat. Mach. Intell.*, vol. 1, no. 5, pp. 206-215, May 2019. doi: 10.1038/s42256-019-0048-x
- [15] J. Slack, S. Hilgard, E. Jia, S. Singh, and H. Lakkaraju, "Fooling LIME and SHAP: Adversarial attacks on post hoc explanation methods", in *Proc. AAAI/ACM Conf. AI Ethics Soc.*, New York, NY, Feb. 2020, pp. 180-186. doi: 10.1145/3375627.3375830
- [16] S. H. Alowais, S. S. Alghamdi, N. Alsuhbany, T. Alqahtani, A. I. Alshaya, S. N. Almohareb, et al., "Revolutionizing healthcare: The role of artificial intelligence in clinical practice", *BMC Med. Educ.*, vol. 23, no. 1, pp. 689, Sep. 2023. doi: 10.1186/s12909-023-04698-z
- [17] D. Alvarez-Melis and T. S. Jaakkola, "On the robustness of interpretability methods," in *Proc. ICML Workshop on Human*

- Interpretability in Machine Learning (WHI 2018)*, Stockholm, Sweden, Jul. 2018, pp. 66–71
- [18] S. Fuchs, B. Olberg, D. Panteli, M. Perleth, and R. Busse, “HTA of medical devices: Challenges and ideas for the future from a European perspective,” *Health Policy*, vol. 121, no. 3, pp. 215–229, 2017, doi: 10.1016/j.healthpol.2016.08.010
- [19] S. Food and Drug Administration, “Artificial Intelligence and Machine Learning (AI/ML)-Enabled Medical Devices”, *FDA Guidance Document*, Silver Spring, MD, Oct. 2023. Available: <https://www.fda.gov/medical-devices/software-medical-device-samd/artificial-intelligence-and-machine-learning-aiml-enabled-medical-devices>
- [20] I. Y. Chen, E. Pierson, S. Rose, S. Joshi, K. Ferryman, and M. Ghassemi, “Ethical machine learning in healthcare”, *Annu. Rev. Biomed. Data Sci.*, vol. 4, pp. 123-144, Jul. 2021. doi: 10.1146/annurev-biodatasci-092820-114757