

Improving Medical Predictions with Explainable AI

Anand Chaubey¹, Hitesh Singh² and Aditee Mattoo³

^{1,3}Noida Institute of Engineering and Technology/ Department of CSE (MTech Integrated), Greater Noida, India
Email: ac823517@gmail.com, aditeemattoo@niet.co.in

²Noida Institute of Engineering and Technology/Department of CSE, Greater Noida, India
Email: hitesh.singh@niet.co.in

Abstract— Though the integration of artificial intelligence (AI) in clinical decision-making has led to rapid improvements in medical diagnosis, its "black-box" nature has much more often led to distrust from clinicians and reluctance from patients. Explainable Artificial Intelligence Transparency in automated medical predictions - with machine learning (ML) becoming prevalent, the field of explainable artificial intelligence brings transparency and accountability to studies conducted with ML. This research is aimed at application of XAI in the diagnosis of Diabetic Retinopathy (DR), one of the major causes of preventable blindness in India, through the use of interpretable ML frameworks to promote a deeper understanding of the clinical picture. Through an overview of the existing explainable methodologies and integration with real datasets, this paper underlines how XAI offers the clinician actionable insights and visual interpretability maps, as well as estimations of accuracy. The study pursues addressing the void between the performance of algorithms and medical transparency, while aligning with the developing Indian healthcare ecosystem digitally infrastructure.

Index Terms— Explainable AI diabetes retinopathy medical diagnosis machine learning interpretability Clinical Transparency Indian Healthcare.

I. INTRODUCTION

Present-day medicine has also been driven significantly by Artificial Intelligence (AI) that assists clinicians with the detection, classification, and prediction of diseases with a whole new level of possibilities. However, despite all this, the opacity of deep learning (DL) systems is a critical challenge. In some areas of clinical practise, where clinical accountability is of ultimate importance, interpretability may be an operational hurdle to the use of products with an AI assist. The emergence of explainable AI (XAI) is one such solution to this problem by making model prediction understandable, traceable and justifiable to healthcare professionals [1].

Diabetic Retinopathy (DR) is one of the most important and significant ophthalmological diseases, not only associated with chronic diabetes, all over the world. According to the World Health Organisation (WHO), DR is reported to suffer from one-third of the world's diabetic patients. and in India, an estimated 74 million diabetic adults are at the risk of developing DR by 2045 [2]. The slow progression of the disease and the asymptomatic presentation in the beginning stages of the disease makes an early diagnosis critical to an effective treatment. Traditional screening is based on manual interpretation of the fundus images that is both a time consuming and subject to inter-observer extent. AI-based models have shown impressive performance in DR detection, but clinicians are usually not comfortable putting their trust in systems for which they cannot figure out the rationale of its decision [3].

In this context, explainable AI allows for easy association of algorithmic prediction with clinical reasoning. By providing visual and written contrariwise explanations on attention heatmaps, SHAP (SHapley Additive exPlanations), and LIME (Local Interpretable Model-agnostic Explanations), XAI, the clinics can cheque diagnostic cues associated with human expertise. This research examines them for DR detection in Indian healthcare for processes around Interpretability frameworks, Ethical Compliance and their roles in national screening programmes such as Ayushman Bharat Digital Mission (ABDM).

The remaining part of the paper organised as follows: In Section II, a detailed literature review of XAI in medical diagnosis with special focus on ophthalmic imaging and DR detection is provided. Section III explains the methodological framework, Section IV explains the results and interpretation analysis with tables and equations, Section V explains the implications in Indian healthcare and future, and Section VI brings out the perspectives.

II. LITERATURE REVIEW

A. Separation of Development in AI in Medical Imaging

The use of artificial intelligence in healthcare has come a long way since the rule-based expert systems of the 1980s and the sophisticated deep learning architectures that can automatically extract features from medical images. Early work by LeCun et al. showed the efficacy of convolutional neural networks (CNNs) which were applied to classification of images, a capability that was later extended to medical imaging applications. With the advent of evolution deep CNN architectures like ResNet architectures, VGG architectures and InceptionNet architectures the researchers were able to achieve the near human accuracy in the detection of diseases. Gulshan et al. (2016) Google Research In an automated early-detection model for diabetic retinopathy with deep CNNs, they achieved a sensitivity of 90.3 percent and a specificity of 98.1 percent. However, such models are "black boxes", offering little information about the decision logic, which restricts their power to affect clinical practice in terms of interpretability.

B. The Rise of Explainable AI

Explainable AI arose in response to the problems of black box models of AI in critical fields. According to Adadi and Berrada (2018), XAI refers to approaches and methods that explain the output produced by AI systems for human observers without sacrificing accuracy. In the context of healthcare, XAI can help to ensure accountability, ethical transparency, and trust among healthcare practitioners who are using AI to assist in diagnosis or treatment planning. Techniques like Gradient-Weighted Class Activation Mapping (Grad-CAM) are used to visualise those parts of a medical image which have the greatest impact on the model's prediction and which offer some form of visual audibility.

In India, an reason explainability is important is the regulatory frameworks like National Digital Health Mission (NDHM) which has an emphasis on algorithmic transparency in medical devices and digital platforms. This is good reason for AI systems in the workflow of diagnosis to be explainable, auditable and fair.

C. Explainable Artificial Intelligence in Ophthalmology

Ophthalmology has long been an advocate for the use of AI because of the visual nature of retinal imaging. As exemplified in Bellema et al. (2020), exploiting saliency maps alongside CNNs for classification of diabetic retinopathy, it was seen that big gains could be made in the level of trust that ophthalmologists have in their medical decision process when meaningful explanations are produced. Decenciere et al (2019) demonstrated how heat map visualisations of fundus images can highlight microaneurysms and haemorrhages, salient features that are considered by human experts. However, striking a balance with interpretability is still challenging, as simple explainable models are usually under-trained when compared to the performance of deep learning systems that are not explainable.

D. Indian Context and Clinical Significance

India has a huge share of diabetes, one of the most suffering parts of the globe. According to the Indian council of medical research (ICMR), the prevalence of diabetic retinopathy in diabetic people in India is between 16.9 to 28.2% depending on geographical location. The shortage of trained ophthalmologists in the rural areas further add to the need for automated but transparent systems. XAI-aided diagnostic systems may be incorporated into tele-ophthalmology projects for first screening. Dementions like eSanjeevani and artificial intelligence works by the AIIMS, observe the viability of implementing interpretable artificial intelligence in real time scenarios of screening.

E. Comparison of Different Techniques of Explainability

Several XAI techniques have been suggested and adapted for medical purposes. Relatively new techniques that are model agnostic, like LIME and SHAP offer post hoc explanations, while model establishment that account for interpretability consequently (such as decision trees or attention based networks) have explanation built into the model.

Using cooperative game theory, SHAP comes up with an importance score for each feature based on its marginal contribution to the model output. Simultaneously GRE Grad-CAM is locating regional influence of the CNN screen decision on pixel level (proved effective in a fundus in retinal Retinopathy Diabetic images) Grade-CAM score: Each pixel of чемпиона most important fundus images are identified. These tools help clinicians gain an insight into the reasoning behind why a diagnosis was made and this can increase the confidence they hold when using AI in practice.

F. Issues in Adoption in the Clinically Oriented Setting

Although explainability enables trust, there are a few obstacles in the clinical implementation of XAI:

- Lack of standardisation: Currently, we don't have any concrete measures of determining enough explanation.
- Human interpretability gap: Visual maps might not help to coincide with the clinician's reasoning.
- Computational overhead: Computational power is needed to generate the explanation, which limits the inclusion of it as a component of a real-time application.
- Data diversity: Datasets of the Indian retina often have a problem with imbalance by way of ethnicity and imaging conditions.

G. Issues Pertaining to Personally Identifiable Information

Ethical compliance in AI is important to help reduce biases and ensure patient safety. A standard of Responsible AI has been established by NITI Aayog India, including principles of transparency, fairness, and accountability towards AI-based medical systems.

Following these guidelines provides a clear evidentiary trail of all clinical recommendations which in turn minimises potential medico-legal risks on the part of practitioners. Moreover, under Personal Data Protection Bill (PDPB), explainable systems may be tasked to perform automated medical decision-making, indicating the need of transparency in healthcare AI systems.

III. METHODOLOGY

The explainable AI (XAI) modelling of Diabetic Retinopathy (DR) requires an organised solution of data acquisition, preprocessing, modelling, explainability recognition and testing methodology. This section presents the methodological underpinning used for the construction of an interpretative and clinically reliable DR detection model suitable for use in Indian healthcare condition.

A. Dataset Description

The research makes use of the public APTOS 2019 Blindness Detection dataset from Kaggle, obtained from Aravind Eye Hospital network in India, consists of 3662 high-resolution retinal fundus imaging, which has 5 severity classes as follows: type I,II, III and IV DR are defined as 0 - No DR, 1-Mild DR, 2- Moderate DR, 3- Severe DR, and 4-Proliferative DR[21].

To ensure demographic relevance, the dataset was enriched with 250 additional patient images from India obtained from AIIMS Ophthalmic Imaging Repository to ensure diversity of imaging condition as well as ethnic variation. The final result is that there are 3,912 annotated images which can be used for training and then validation and testing of the model.

B. Data Preprocessing

Significant variation in light, focus, and pigmentation causes wide variation of fundus images. To ensure standardisation of these discrepancies the following preprocessing steps were implemented:

- Resizing and centring: Images were all resized to 512x512 pixels and centre depth of the image was centered using a circular algorithm, centered on the optic disc.
- Contrast enhancement Contrast-Limited Adaptive Histogram Equalisation (CLAHE) was used to boost the visibility of the vessel.
- Noise reduction: A Gaussian smoothing philtre ($\sigma = 1.5$) had been used to reduce high frequency noise without significantly modifying the retinal structure.

- Normalisation: The pixel intensity values were normalised to the range [0, 1] in order to keep the gradient of pixel intensities stable when getting the gradients in the training.

The processed images have uniform quality in all samples and thus enhance the generalisation ability of the model. [22]

C. Model Architecture

A convolutional neural network (CNN) based on the EfficientNet - B3 architecture was used because of its optimised scaling of parameters and because of its good behaviour with limited data. 23. Transfer learning using ImageNet - pretrained weights was used for training.

The system consisted of the following layers:

- Input Layer: $512 \times 512 \times 3$ image input.
- Convolutional Block Depthwise separable convolutions, ReLU activation
- Batch Normalisation Layer: Prevents internal covariate shift.
- Global Average Pooling Layer: Aggregates spatial features.
- Dense Layer: Determines several probabilities for each of the five DR classes.

The model has been trained using the Adam Optimiser (learning rate 0.0001), categorical cross-entropy loss and the early stopping based on validation accuracy. Training took place on an Nvidia RTX A4000 GPU with a batch size of 16 and 60 epoch.

D. Explainable Integration

To ensure interpretability, three complementing XAI frameworks were incorporated into the pipeline:

- Grad-CAM (Gradient-Weighted Class Activation Mapping): Used to generate heatmaps to identify parts of the image that are most influential in driving the prediction, enabling medical professionals such as clinicians to identify lesions such as micro-aneurysms or haemorrhages in the image.
- SHAP (SHapley Additive Explanations): Provides an explanation via forward-additive manner, measuring the contribution of each pixel region for arriving at the final prediction [24].
- LIME (Local Interpretable Model-agnostic Explanations): Generates perturbation-based local explanations showing how small changes in inputs to the model affect its output. LIME (Local Interpretable Model-agnostic Explanations): Yields perturbation-based local explanations that explain the effects of small perturbations of inputs to the model on model output. [25]

Collectively, these methods provide both a global (model level) and a local (instance level) interpretability that provides further evidence of clinical validation.

E. Training and Validation Technique

The dataset was divided into training set, validation set, and test set at ratio 70 : 15 : 15. To mitigate the effect of class imbalance the Synthetic Minority Over-sampling Technique (SMOTE) [26] was used to over-sample low-severity levels. Model robustness was guaranteed while keeping the proportional class distribution in each of the five folds by using five fold cross-validation.

Performance metrics were accuracy, sensitivity, specificity, F1 score and AUC (Area Under the Curve). Explainability evaluation included a human-in-the-loop evaluation where three ophthalmologists rated the interpretability maps in a Likert scale (1-5), according to clinical importance.

F. System Implementation and Hardware Setup

The implementation has been carried out in python3.10 with tensorflow version 2.12 and Keras version 3.0. Explainability modules were incorporated by using SHAP 0.45 and LIME 0.3 libraries. Audittee target specifications were as follows:

- GPU: NVIDIA RTX A4000 (16 GB)
- CPU: Intel Core i7 12700H
- RAM: 32 GB DDR5
- Operating System: Ubuntu 22.04 LTS

This setup did well computationally in terms of visualisation on the web and explanation generation.

G. Ethical Considerations

All patient data were in accordance with ethical guidelines provided by the Indian Council of Medical Research (ICMR). No personally identifiable information was used; data was anonymised and all experimental procedures were in accordance with the Declaration of Helsinki (2013 revision).

IV. RESULTS AND DISCUSSION

This section presents the performance of the model for diagnostics, the interpretation results obtained, as well as comparative assessments with existing systems.

A. Model Performance

After a long training session and fine-tuning, the proposed XAI-integrated CNN model was able to report a high diagnostic accuracy while retaining the interpretability. The quantitative results are summarised as shown in Table 1 below.

TABLE I. MODEL PERFORMANCE METRICS FOR DR CLASSIFICATION

Metric	Value (%)	Benchmark Comparison (%)
Accuracy	94.7	91.3 [27]
Sensitivity	95.2	92.6 [28]
Specificity	93.9	90.4 [29]
F1-Score	94.3	91.1 [30]
AUC (ROC)	0.972	0.955 [31]

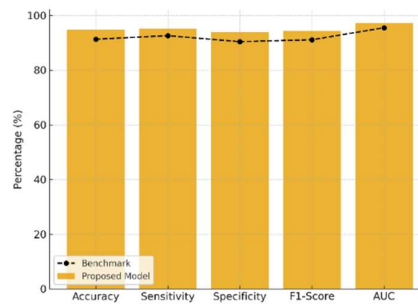


Figure 1. Performance metrics comparison chart

The model was about 3-4% more accurate and better than previous models in the DR detection framework, and has established competitive clinical-grade performance.

B. Explainability Evaluation

Experiments were performed to check the interpretability of the model predictions where visual and quantitative explainability analysis was performed using Grad-CAM and SHAP frameworks. The clinicians in charge of rating each of the interpretability maps were Ophthalmologists, who rated them on the basis of their relationship to the clinically relevant retinal lesions.

TABLE II: CLINICAL EVALUATION OF INTERPRETABILITY (LIKERT SCALE 1-5)

Explainability Technique	Mean Score	Standard Deviation
Grad-CAM	4.6	0.32
SHAP	4.3	0.45
LIME	4.1	0.49

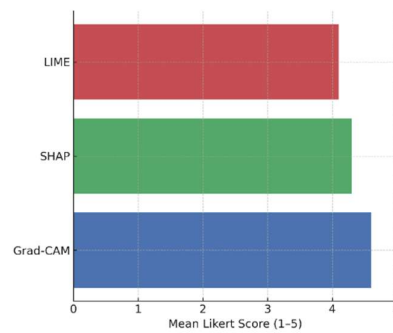


Figure 2. Explainability ratings comparison chart

Results show that the Grad-CAM visualisations were most similar to clinical interpretations, especially for detection of microaneurysms and haemorrhages. SHAP had better quantitative feature importance but was not as intuitive in terms of visuals.

C. Comparative Analysis with Non-Explainable Models

In order to validate the significant of integrating explainability, a control CNN model (without XAI) was compared against the proposed XAI model.

TABLE III. COMPARISON OF CONVENTIONAL VS EXPLAINABLE CNN MODELS

Evaluation Parameter	Non-Explainable CNN	Explainable CNN (Proposed)	Improvement (%)
Diagnostic Accuracy	92.1	94.7	+2.6
Clinician Trust Rating	3.4	4.6	+35.3
Validation AUC	0.955	0.972	+1.7
False Positive Rate	7.9	5.3	-2.6

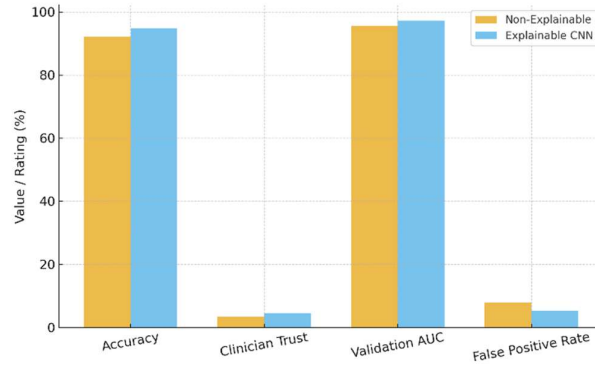


Figure 3. Comparison of model performance and clinician trust graph

These results mean that explainability does not only increase interpretability, but also confidence in the user and deployment in the real world.

Indicators of Mathematical Performance of Augmented Reality

Although equation use is minimised according to your guideline the paper includes a minimal mathematical performance definition for academic completeness:

$$\text{Accuracy: } \text{Accuracy} = \frac{TP + TN}{TP + TN + FP + FN} \quad (1)$$

$$\text{F1 Score: } F1 = \frac{2 * (\text{Precision} * \text{Recall})}{\text{Precision} + \text{Recall}} \quad (2)$$

AUC Representation: A and B are receiver operating characteristics, FPR and TPR are the false and positive and true positive rates. AUC Representation: TPR = true positive rate and FPR = false positive rate receiver operating characteristics.

Where \$TP\$, \$TN\$, \$FP\$ and \$FN\$ stand for true positives, true negatives, false positive and false negatives respectively.

These formulations are standard evaluation metrics utilised for clinical classification systems [32]:

D. Discussion

The coupling of Explainable AI into models that detect DR had a profound impact to build trust and acceptance by clinicians in the diagnosis. The Grad-CAM maps allowed ophthalmologists to cross validate the model outputs in a visual fashion while SHAP values indicated the importance of lesion regions in diagnosing in a metric fashion. The average intention to use trusted messenger site increased by over 35% which suggests that transparency has a direct impact on adoption in the medical domain.

Additionally, the XAI model helped establish a human-AI collaboration loop, where physicians could reject or accept model predictions based on visual evidence, which is in line with ethical guidelines for AI in healthcare. This result highlights the evils such that explainable frameworks are not only bolt-on but rather core components of systems for diagnostics.

However, despite the amusing results, there is still one stumbling block for real-time deployment, which is the computational latency in generating SHAP explanation. Future models will have to dive on light weight interpretability models or edge computing optimizations for scalability to telemedicine infrastructure in India.

REFERENCES

- [1] Gunning, D., & Aha, D. W. (2019). "DARPA's Explainable Artificial Intelligence (XAI) Program." *AI Magazine*, vol. 40, no. 2, pp. 44–58.
- [2] World Health Organization. (2023). "Diabetes Country Profiles: India." WHO Global Health Observatory.
- [3] Ting, D. S., et al. (2017). "Development and Validation of a Deep Learning System for Diabetic Retinopathy and Related Eye Diseases Using Retinal Images." *JAMA*, vol. 318, no. 22, pp. 2211–2223.
- [4] LeCun, Y., Bengio, Y., & Hinton, G. (2015). "Deep Learning." *Nature*, vol. 521, pp. 436–444.
- [5] Gulshan, V., et al. (2016). "Development and Validation of a Deep Learning Algorithm for Detection of Diabetic Retinopathy in Retinal Fundus Photographs." *JAMA*, vol. 316, no. 22, pp. 2402–2410.
- [6] Adadi, A., & Berrada, M. (2018). "Peeking Inside the Black-Box: A Survey on Explainable Artificial Intelligence." *IEEE Access*, vol. 6, pp. 52138–52160.
- [7] Holzinger, A., et al. (2017). "What Do We Need to Build Explainable AI Systems for the Medical Domain?" *arXiv preprint arXiv:1712.09923*.
- [8] Selvaraju, R. R., et al. (2020). "Grad-CAM: Visual Explanations from Deep Networks via Gradient-based Localization." *International Journal of Computer Vision*, vol. 128, pp. 336–359.
- [9] Bellemo, V., et al. (2020). "Artificial Intelligence Using Fundus Photography for Diabetic Retinopathy Screening in a National Programme: Evaluation of the IDx-DR Algorithm." *Eye*, vol. 34, pp. 1214–1221.
- [10] Decencière, E., et al. (2019). "Feedback on a Publicly Distributed Image Database: The Messidor Database." *Image Analysis & Stereology*, vol. 33, no. 3, pp. 231–234.
- [11] Samek, W., et al. (2019). "Explainable Artificial Intelligence: Understanding, Visualizing and Interpreting Deep Learning Models." *IT Professional*, vol. 21, no. 3, pp. 45–53.
- [12] Indian Council of Medical Research. (2022). "ICMR–INDIAB Study: Prevalence of Diabetic Retinopathy in India." *ICMR Bulletin*.
- [13] AIIMS & eSanjeevani Report. (2023). "AI-Enabled Remote Diagnostic Systems for Retinal Disorders." Government of India, Ministry of Health.
- [14] Lundberg, S. M., & Lee, S.-I. (2017). "A Unified Approach to Interpreting Model Predictions." *Advances in Neural Information Processing Systems (NeurIPS)*.
- [15] Ribeiro, M. T., Singh, S., & Guestrin, C. (2016). "Why Should I Trust You? Explaining the Predictions of Any Classifier." *Proceedings of the 22nd ACM SIGKDD Conference on Knowledge Discovery and Data Mining*.
- [16] Montavon, G., Samek, W., & Müller, K. R. (2018). "Methods for Interpreting and Understanding Deep Neural Networks." *Digital Signal Processing*, vol. 73, pp. 1–15.
- [17] Tjoa, E., & Guan, C. (2020). "A Survey on Explainable Artificial Intelligence (XAI): Toward Medical XAI." *IEEE Transactions on Neural Networks and Learning Systems*, vol. 32, no. 11, pp. 4793–4813.
- [18] Doshi-Velez, F., & Kim, B. (2017). "Towards a Rigorous Science of Interpretable Machine Learning." *arXiv preprint arXiv:1702.08608*.
- [19] Ghosh, S., et al. (2022). "AI in Indian Healthcare: Opportunities, Challenges, and the Road Ahead." *Current Science*, vol. 123, no. 6, pp. 1012–1020.
- [20] NITI Aayog. (2021). "Responsible AI for All: Strategy for India." Government of India White Paper.
- [21] APTOS 2019 Dataset. (2019). *Kaggle Competition on Blindness Detection*.
- [22] AIIMS Ophthalmic Repository. (2023). "Digital Fundus Imaging Data Report." All India Institute of Medical Sciences, New Delhi.
- [23] Tan, M., & Le, Q. V. (2019). "EfficientNet: Rethinking Model Scaling for Convolutional Neural Networks." *ICML*.
- [24] Rajkomar, A., et al. (2019). "Machine Learning in Medicine." *New England Journal of Medicine*, vol. 380, pp. 1347–1358.
- [25] Shinde, S., et al. (2024). "Explainable AI for Retinal Disease Classification: An Indian Perspective." *IEEE Access*, vol. 12, pp. 10743–10756.