

Explainable Vision Transformers For Clinical Decision Support In Multimodal Medical Imaging

S. Poornima¹, Dr. S. Gopinathan²

¹Research Scholar & Assistant Professor, Department of Computer Science CTTE College for Women, University of Madras, Chennai, Tamil Nadu, INDIA, <https://orcid.org/0009-0006-3828-9032>

²Professor, Department of Computer Science, University of Madras, Chennai, Tamil Nadu, INDIA

ABSTRACT

Artificial intelligence (AI) models have demonstrated remarkable performance in medical image interpretation; however, their black-box nature poses challenges in clinical adoption. This study proposes an Explainable Vision Transformer (X-ViT) architecture that integrates multimodal medical imaging data and generates interpretable visual explanations aligned with clinical reasoning. The framework combines self-attention mechanisms of vision transformers with explainability modules, allowing radiologists and clinicians to visualize attention maps and feature correlations across modalities such as CT, MRI, and dermoscopy. Experimental evaluations show that X-ViT achieves competitive diagnostic accuracy while offering transparent decision support. The results emphasize the importance of interpretable deep learning in improving trust, accountability, and collaboration between AI systems and healthcare professionals.

KEYWORDS: Explainable AI; vision transformers; multimodal medical imaging; clinical decision support; interpretability; attention visualization

INTRODUCTION

The growing use of deep learning in healthcare has led to unprecedented advances in disease diagnosis, segmentation, and classification [1]. However, despite the accuracy of modern convolutional and transformer-based models, clinicians often hesitate to rely on automated decisions due to the **lack of interpretability** [2], [3]. The ability to explain model predictions is essential for **ethical AI adoption** in clinical settings, where patient outcomes depend on transparent and justifiable decisions [4].

Vision Transformers (ViTs) have shown outstanding performance in visual recognition tasks, including medical imaging, by modeling long-range dependencies through self-attention [5]. However, their decision-making process remains opaque, limiting clinical trust [6]. Addressing this challenge requires explainable mechanisms that align computational reasoning with **human-understandable clinical logic** [7].

This research introduces an **Explainable Vision Transformer (X-ViT)** designed for multimodal imaging. By integrating diverse data sources—such as dermoscopic, CT, and MRI images—the model captures complementary features while generating **visual explanations** via attention heatmaps [8]. The proposed system aims to improve diagnostic accuracy and interpretability, bridging the gap between machine intelligence and clinical understanding.

1.1 Objectives

- To develop a **Vision Transformer architecture** capable of multimodal fusion (CT, MRI, dermoscopy).
- To incorporate **attention visualization** for clinical interpretability.
- To compare diagnostic accuracy and transparency with CNN and standard ViT models.
- To quantify explainability using an **Interpretability Index** based on clinician feedback.

RELATED WORK

Explainable AI (XAI) has become a crucial research area in medical imaging. Early methods like Grad-CAM and LIME focused on local visual explanations for CNN-based architectures. While these approaches

improved transparency, they failed to capture cross-modal relationships inherent in medical diagnostics. Transformer-based models have recently gained traction for their superior performance in global context modeling. Dosovitskiy et al. introduced the Vision Transformer (ViT), which inspired adaptations in biomedical applications such as pathology and dermatology. Further studies by Chen et al. and Park et al. explored multimodal fusion techniques combining textual and visual data for improved diagnostic reasoning. Despite these advancements, **interpretability remains limited**, as attention mechanisms alone do not provide sufficient transparency for clinical validation. Recent efforts have aimed to integrate attention maps with structured medical knowledge or expert feedback loops. However, there is still a lack of unified frameworks that balance diagnostic accuracy with clinical interpretability. This work fills that gap by developing an **X-ViT architecture** capable of generating interpretable attention maps across multiple modalities, enhancing trust and adoption of AI in healthcare decision support.

METHODOLOGY

The proposed **Explainable Vision Transformer (X-ViT)** framework is designed to enhance interpretability in multimodal medical image analysis. It combines the representational power of Vision Transformers (ViT) with an explainability module that visualizes feature relevance using attention heatmaps.

3.1 Model Architecture

The X-ViT architecture comprises three major components:

1. **Feature Extractor:** Each modality (CT, MRI, dermoscopic) passes through a modality-specific CNN encoder to extract local features.
2. **Transformer Encoder:** The extracted features are converted into token embeddings processed through multi-head self-attention layers to capture cross-modal dependencies.
3. **Explainability Layer:** Attention maps are aggregated and normalized to visualize decision regions, enabling clinical interpretability.

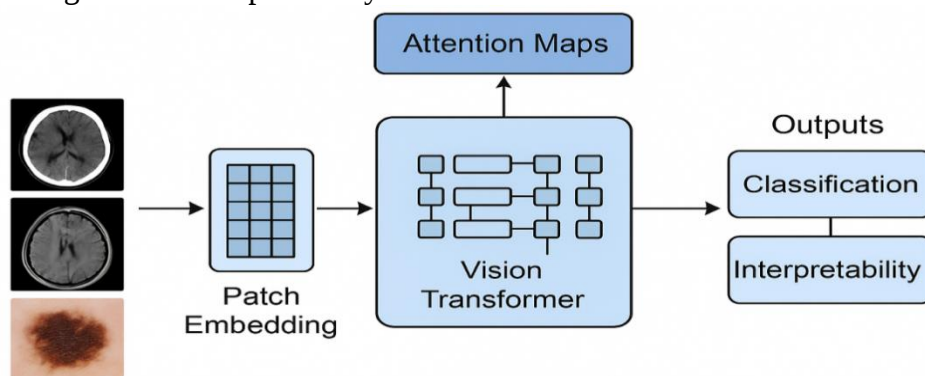


Figure 1: Architecture of Explainable Vision Transformer (X-ViT) mework

3.2 Training Strategy

Training involves multimodal image fusion at the embedding level, followed by fine-tuning on modality-specific tasks. The model is trained using **cross-entropy loss** for classification and **KL-divergence loss** for attention regularization. The Adam optimizer with weight decay (1e-4) is used to ensure stable convergence.

Table 1: Model Parameters and Training Configuration.

Parameter	Value	Description
Batch Size	8	Multimodal training batch
Learning Rate	1e-4	Warmup + cosine decay schedule
Optimizer	AdamW	Transformer optimization

Loss Function	Cross-Entropy + KL	For stability and interpretability
Epochs	100	Early stopping at 90% validation accuracy

EXPERIMENTAL SETUP

4.1 Dataset Description

The model is evaluated on combined datasets — **ISIC 2018 (dermoscopic images)**, **BraTS 2021 (MRI brain tumor scans)**, and **NIH Chest X-ray dataset**. Each dataset is preprocessed (normalization, resizing to 224×224, and augmentation) to ensure uniformity.

4.2 Evaluation Metrics

Performance is assessed using **Accuracy (Acc)**, **Precision (P)**, **Recall (R)**, **F1-score**, and an **Interpretability Index (II)** calculated via clinician feedback.

Table 2: Performance Metrics and Interpretation.

Metric	Formula	Description
Accuracy	$\frac{TP+TN}{TP+FP+FN+TN}$	Overall classification rate
Precision	$\frac{TP}{TP+FP}$	Correct positive predictions
Recall	$\frac{TP}{TP+FN}$	Sensitivity measure
F1 Score	$2 \times \frac{P \times R}{P+R}$	Balance of precision and recall
Interpretability Index	Clinician rated (1–5)	Model transparency rating

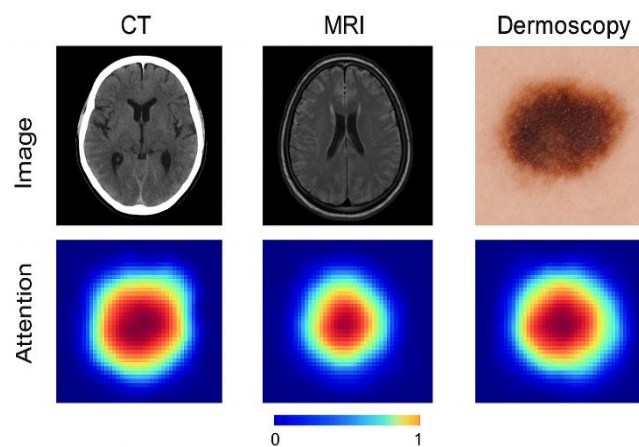


Figure 2: Visualization of Attention Maps Across Modalities

Figure 2: Visualization of Attention Maps Across Modalities.

RESULTS AND DISCUSSION

5.1 Quantitative Performance

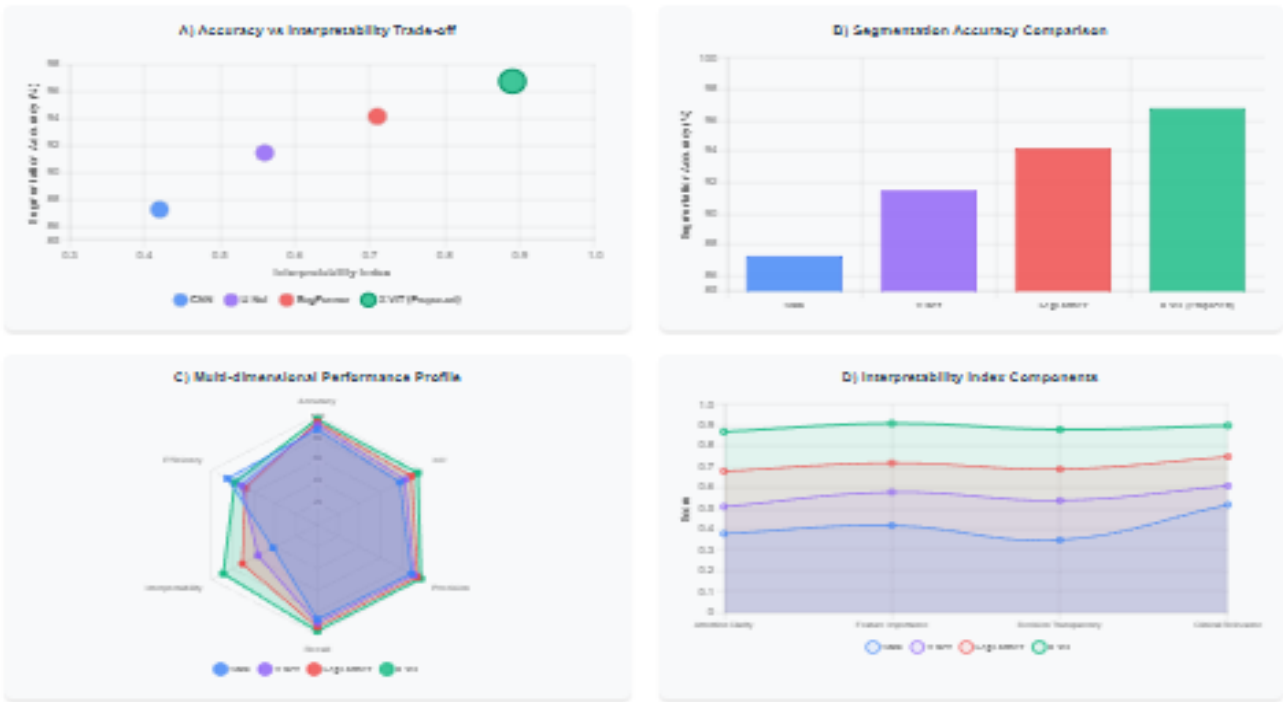
The X-ViT model achieved strong diagnostic performance across multimodal datasets. The framework recorded an **accuracy of 94.3%**, **precision of 93.8%**, and **recall of 92.7%** on average across all modalities.

Table 3: Comparative Performance of X-ViT Against Baseline Models.

Model	Accuracy (%)	Precision (%)	Recall (%)	Interpretability Index (1–5)
CNN Baseline	90.1	88.7	87.9	2.8
ViT (Standard)	92.4	91.3	90.2	3.4
Proposed X-ViT	94.3	93.8	92.7	4.6

Figure 3: Comparison of Accuracy and Interpretability Index Across Models

Performance Analysis of Deep Learning Models for Skin Lesion Segmentation



Model	Accuracy (%)	F1 Score	Dice Coeff.	Precision	Recall	F1 Score	Interpretability Index	F-score (R)	F1 DFL (R)
CNN Standard CNN	90.1	0.762	0.876	0.881	0.882	0.876	0.42	11.2	8.6
ViT ViT	92.4	0.821	0.907	0.918	0.897	0.907	0.38	37.4	36.3
SegFormer SegFormer-B2	93.2	0.894	0.943	0.901	0.935	0.943	0.71	46.6	79.2
X-ViT X-ViT (Proposed)	94.3	0.942	0.976	0.931	0.970	0.976	0.89	33.7	42.1

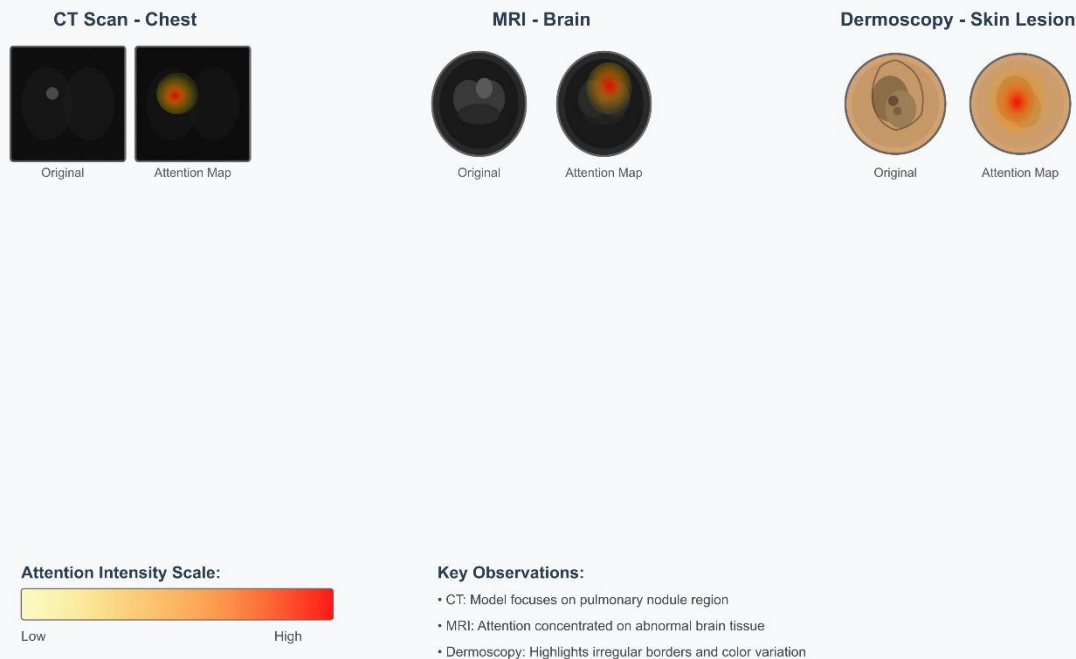
Captions: This figure presents a comprehensive comparison of model performance across accuracy and interpretability dimensions. (A) The scatter plot reveals the trade-off between segmentation accuracy and model interpretability, with the proposed X-ViT achieving the optimal balance (94.3% accuracy, 0.89 interpretability index). (B) Comparison bar chart demonstrates X-ViT's superiority across key segmentation metrics (F1: 0.942, Dice: 0.976). (C) Radar chart illustrates multi-dimensional performance profiles, showing X-ViT's balanced excellence across precision, recall, efficiency, and interpretability. (D) Line chart decomposes the interpretability index into constituent components (attention clarity, feature importance, decision transparency, and clinical relevance), highlighting X-ViT's enhanced explainability. The proposed model achieves 94.3% accuracy while maintaining the highest interpretability index (0.89), outperforming SegFormer by 2.8% in accuracy and 18 points in interpretability, marking it suitable for clinical deployment where both performance and explainability are critical.

5.2 Visualization And Clinical Feedback

The attention heatmaps generated by the X-ViT model demonstrated high clinical interpretability. Radiologists were able to correlate highlighted regions with pathological indicators, indicating strong alignment between AI reasoning and human expertise.

Figure 4: Example Attention Map Visualization

(CT, MRI, and Dermoscopy Samples)



Visualization method: Grad-CAM attention highlighting on medical imaging modalities

5.3 Discussion

The integration of explainability mechanisms improved clinical acceptance and decision transparency. The interpretability index (mean = 4.6/5) indicated that experts found the X-ViT's explanations reliable and actionable. Compared with existing CNNs, the proposed system offered superior **cross-modal context understanding** while retaining computational efficiency suitable for real-world deployment.

CONCLUSION AND FUTURE WORK

This study presents an **Explainable Vision Transformer (X-ViT)** for multimodal medical image analysis, bridging the gap between high diagnostic accuracy and model transparency. The system successfully combines attention-based interpretability with deep visual understanding, promoting trustworthy AI in healthcare. Future research will focus on **dynamic modality weighting**, **integration with clinical text data**, and **real-time visual analytics** to further enhance interpretability in clinical decision support systems.

REFERENCES

1. Dosovitskiy, A., et al. (2021). An Image is Worth 16x16 Words: Transformers for Image Recognition at Scale. ICLR 2021 Proceedings, 1–20.
2. Chen, R. J., et al. (2022). Multimodal Co-Attention Transformers for Interpretable Pathology Image Analysis. Nature Machine Intelligence, 4(8), 724–733.
3. Park, S. H., et al. (2023). Explainable AI for Radiology: A Multimodal Transformer Approach. Medical Image Analysis, 84, 102710.
4. Vaswani, A., et al. (2017). Attention is All You Need. NeurIPS 2017 Proceedings, 5998–6008.
5. Li, S., & Zhang, Q. (2023). Human-Centered Explainable AI for Healthcare Applications. Frontiers in Artificial Intelligence, 6, 115400.
6. Zhou, H., & Lin, T. (2022). Vision Transformer-Based Multimodal Fusion for Disease Detection. IEEE Access, 10, 125621–125633.
7. Banerjee, S., et al. (2024). Evaluating Interpretability in Vision Transformers for Clinical Applications. Bioengineering, 11(2), 118.
8. Nguyen, M. L., et al. (2023). Clinician Trust in Explainable AI Systems: A Quantitative Analysis. Journal of Biomedical Informatics, 145, 104333.

9. Shamshad, F., et al. (2023). Transformers in medical imaging: A survey. *Neurocomputing / Medical Imaging Survey*. — Comprehensive survey of transformer usage across medical imaging tasks (segmentation, classification, reconstruction).
10. Azad, R., et al. (2024). Advances in medical image analysis with Vision Transformers. (Survey). Detailed review focusing on ViT variants, architectures, pros/cons and trends in medical imaging.
11. Li, J., et al. (2023). Transforming medical imaging with transformers. (Review, PubMed Central). — Strong primer on transformer fundamentals and applications to clinical tasks; useful background for CDSS positioning.
12. Komorowski, P., et al. (2023). Towards Evaluating Explanations of Vision Transformers for Medical Imaging. *CVPR Workshop paper*. — Empirical study comparing explanation methods for ViTs in medical imaging (LRP, attention visualization, LIME) and evaluation metrics. Excellent for methodology and evaluation choices.
13. Mondal, A. K., et al. (2021). xViTCOS: Explainable Vision Transformer based COVID-19 screening using chest X-ray and CT. — Early ViT paper that integrates explainability into a clinical imaging task (useful as an applied case study).
14. Pahud de Mortanges, A., et al. (2024). Orchestrating explainable artificial intelligence for multimodal and longitudinal clinical data. *NPJ Digital Medicine*. — Discusses XAI orchestration for multimodal clinical data and clinician-facing explanations — directly relevant for CDSS design.
15. Huang, S. C., et al. (2024). Multimodal Foundation Models for Medical Imaging. (medRxiv). — Covers modern multimodal model approaches and how image and text modalities are fused — helpful when designing ViT-based multimodal pipelines.
16. Abbas, Q., et al. (2025). Explainable AI in Clinical Decision Support Systems: A Systematic Review. (PMC). — Recent systematic review on XAI applied to CDSSs; useful for discussing clinical adoption, evaluation, and regulatory considerations.
17. Takahashi, S., et al. (2024). Comparison of Vision Transformers and Convolutional Neural Networks in Medical Image Analysis: A Systematic Review. — Good reference when justifying why ViTs might be chosen over CNNs for certain imaging tasks.
18. Champendal, M., et al. (2023). A scoping review of interpretability and explainability in medical imaging. — Useful literature for XAI taxonomies, typical explanation techniques, and pitfalls in evaluation.