

# ViT-WoundNet: An Explainable Deep Learning Framework for Wound Type and Severity Classification

Rupesh Kumar Sah<sup>✉</sup>

Computer Science and Engineering  
Indian Institute of Technology, Patna  
Bihar, India  
rupesh\_2421cs03@iitp.ac.in

Anil Verma<sup>✉</sup>

Computer Science and Engineering  
Indian Institute of Technology, Patna  
Bihar, India  
anil\_2321cs11@iitp.ac.in

Kumar Prasun<sup>✉</sup>

Computer Science and Engineering  
Indian Institute of Technology, Patna  
Bihar, India  
kumar\_2422cs12@iitp.ac.in

Rajiv Misra<sup>✉</sup>

Computer Science and Engineering  
Indian Institute of Technology, Patna  
Bihar, India  
rajivm@iitp.ac.in

**Abstract**—Deep learning advances have revolutionized automated wound assessment, yet convolutional neural networks (CNNs) struggle with inter-class similarities, illumination variations, and generalization across diverse datasets. This work introduces ViT-WoundNet, a Vision Transformer (ViT)-based framework leveraging self-attention to capture long-range spatial dependencies and global context, addressing these limitations effectively. ViT-WoundNet outperforms CNN baselines—ConvNeXt (85.4%), EfficientNet (86.7%), and ResNet50 (74.1%)—achieving 91.3% accuracy for wound-type classification and 88.4% for severity across 2,372 multi-class images from five public datasets. Confusion matrices and Grad-CAM visualizations confirm superior inter-class separability and clinically interpretable feature localization, positioning ViT-WoundNet as a robust, explainable solution for telemedicine applications.

**Index Terms**—Vision Transformer, Wound Classification, Medical Image Analysis, Explainable AI, Deep Learning.

## I. INTRODUCTION

Recent advances in deep learning and computer vision have revolutionized medical image analysis by enabling end-to-end feature learning from raw pixels, bypassing handcrafted features [1]. CNNs and ViT-WoundNet excelled across dermatology, ophthalmology, radiology, and wound assessment [2].

Automated wound analysis nevertheless faces critical barriers to clinical deployment. Most studies employed small proprietary datasets for binary tasks, neglecting diverse multi-class wound morphologies [1]. CNNs effectively captured local textures but struggled with irregular boundaries, inter-class similarities (e.g., venous vs. diabetic ulcers), clinical photo illumination variations, and long-range spatial dependencies constrained by fixed receptive fields.

**Research Gap:** State-of-the-art CNNs such as ResNet [3], EfficientNet [4], and ConvNeXt nevertheless exhibit architectural limitations for wound analysis. Comprehensive multi-class datasets encompassing 10+ wound types across 3 sever-

ity grades remain scarce, while explainability—critical for clinicians—is underexplored; visualizations must align with protocols for necrotic tissue localization and granulation identification. ViTs provide superior global reasoning via self-attention [5], yet their use for joint wound type/severity classification sees limited adoption.

This work addresses these gaps through four key contributions:

(1) **Multi-source Dataset:** Five public wound datasets are consolidated (2,372 images total) across 10 wound types and 3 severity levels, yielding the largest open-access multi-class benchmark through stratified sampling.

(2) **ViT-WoundNet Architecture:** A dual-head Vision Transformer is proposed with task-specific heads after the CLS token, trained via balanced multi-task loss  $\mathcal{L} = 0.6\mathcal{L}_{\text{type}} + 0.4\mathcal{L}_{\text{severity}}$ .

(3) **Comprehensive Benchmarking:** ViT-WoundNet attains 91.3% accuracy on wound-type classification and 88.4% on severity, outperforming CNN baselines (ConvNeXt: 85.4%, EfficientNet-B0: 86.7%, ResNet50: 74.1%).

(4) **Clinical Explainability:** Grad-CAM visualizations align attention with necrotic cores and ulcer margins [6], while confusion matrices demonstrate enhanced inter-class separability.

These advances position ViT-WoundNet as a state-of-the-art, interpretable solution for automated wound assessment, with implications for telemedicine and resource-limited settings.

## II. RELATED WORKS

Wound image analysis has evolved from handcrafted features to deep learning for classification, segmentation, and severity assessment [1].

### A. Early CNN Approaches

Ensemble CNNs achieved 85% accuracy for multi-class wound classification on limited custom datasets, outperform-

ing traditional methods but lacking interpretability [2]. Eff-ReLU-Net improved feature extraction on Kaggle datasets (88% accuracy) through specialized activations [7].

### B. Multi-modal and Hybrid Architectures

Multi-modal fusion of wound images with anatomical location data using ResNet152/EfficientNet variants reported 78–100% accuracy across four classes [8]. CNN-Transformer hybrids incorporated attention for better representation but remained limited to small-scale tasks [9].

### C. Specialized Severity Classification

Pressure ulcer staging employed deep CNNs [10], while EfficientNet excelled in diabetic foot ulcer detection [11]. HLG-Net achieved 82.6% accuracy on four wound classes via tissue-specific features [12].

### D. Vision Transformers in Wound Analysis

ViT-WoundNet demonstrated promise for wound stage classification (90% accuracy) through global self-attention, outperforming CNNs on severity tasks [13].

TABLE I: SOTA Wound Classification Methods

| Method           | Dataset | Classes  | Acc. (%) | Explainability |
|------------------|---------|----------|----------|----------------|
| Ensemble CNN [2] | Custom  | Multi    | 85       | None           |
| Eff-ReLU-Net [7] | Kaggle  | Multi    | 88       | None           |
| Multi-modal [8]  | AZH     | 4        | 78–100   | Partial        |
| ViT Stage [13]   | Custom  | Severity | 90       | Grad-CAM       |
| HLG-Net [12]     | Custom  | 4        | 82.6     | None           |

This SOTA review reveals the progression from CNN ensembles to transformer-based methods, with persistent limitations in dataset scale, multi-task capability, and comprehensive explainability.

## III. METHODOLOGY

This section presents dataset compilation, preprocessing with augmentation, and model training with optimization to achieve accurate and interpretable clinical predictions.

### A. Dataset and Preprocessing

To ensure diversity and generalization, five publicly available wound image datasets were combined. The sources include the **Wound Dataset** by Pratomo (Kaggle, 2024), the **Wound Classification Dataset** by Fateen (Kaggle, 2024), the **Wound Segmentation Dataset** by Leoscode (Kaggle, 2024), the **Wound Data** by Taher (Kaggle, 2024), and the **Skin Burn Dataset** by Baid (Kaggle, 2024). Together, these datasets cover a wide range of wound types (*abrasion, bruises, burn, cut, diabetic, laceration, pressure, surgical, venous, normal*) and three severity grades (*mild, moderate, severe*). Images were resized to  $224 \times 224$  and normalized with ImageNet means/standard deviations. The training set used on-the-fly augmentation: random horizontal/vertical flips,  $\pm 15^\circ$  rotations, and color jitter (brightness/contrast/saturation), while validation/test sets used only resizing and normalization. The dataset was split with stratification into 70% train, 15% validation, and 15% test via a two-stage split (30% hold-out then 50/50).

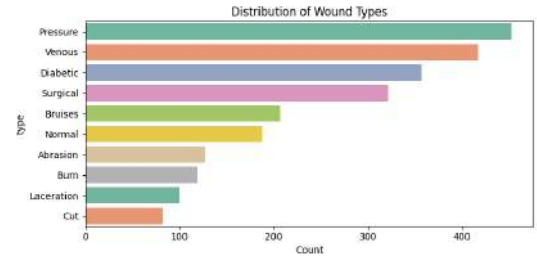


Fig. 1: Distribution of wound types across the combined dataset (abrasion, bruises, burn, cut, diabetic, laceration, pressure, surgical, venous, normal).

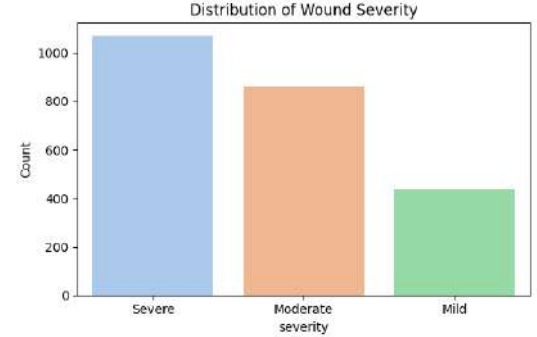


Fig. 2: Distribution of wound severity levels (mild, moderate, severe) across the combined dataset.

Standard preprocessing techniques were used to normalize and enhance the wound images before model training. Each image was resized to  $224 \times 224$  pixels, converted to RGB format, and normalized to the  $[0, 1]$  range using min-max normalization as expressed in (1):

$$I_{norm} = \frac{I - I_{min}}{I_{max} - I_{min}} \quad (1)$$

where  $I_{min}$  and  $I_{max}$  represent the minimum and maximum pixel values, respectively. Data augmentation (rotation  $\pm 15^\circ$ , horizontal/vertical flips, random brightness shifts between 0.8–1.2, and Gaussian noise) was applied to improve generalization. Descriptive statistics were calculated to understand brightness and color variations. The mean ( $\mu$ ) and standard deviation ( $\sigma$ ) of pixel intensities were computed as in (2):

$$\mu = \frac{1}{N} \sum_{i=1}^N I_i, \quad \sigma = \sqrt{\frac{1}{N} \sum_{i=1}^N (I_i - \mu)^2} \quad (2)$$

where  $I_i$  is the intensity of pixel  $i$  and  $N$  is the total number of pixels. These measures helped identify and correct illumination inconsistencies across the dataset.

### B. Model Architectures

Four architectures—ConvNeXt, ResNet-50, EfficientNet-B0, and the proposed (ViT-WoundNet)—were implemented

and compared for multi-class wound-type and severity classification. All models were initialized with pretrained ImageNet-1K weights and fine-tuned using the AdamW optimizer with a learning rate of  $1 \times 10^{-4}$ , batch size of 16, and early stopping based on validation loss. Each model was trained for 50 epochs with categorical cross-entropy loss and weight decay of  $1 \times 10^{-5}$  to prevent overfitting. The hyperparameter settings summarized in Table II.

TABLE II: Hyperparameters used in all experiments

| Hyperparameter                     | Value                                                               |
|------------------------------------|---------------------------------------------------------------------|
| Learning rate                      | $9.14548 \times 10^{-6}$                                            |
| Batch size                         | 32                                                                  |
| Fine-tune last $N$ backbone blocks | 2                                                                   |
| Dropout (per output head)          | 0.5                                                                 |
| Max epochs                         | 50                                                                  |
| Optimizer                          | Adam (wd = $1 \times 10^{-4}$ )                                     |
| LR scheduler                       | ReduceLROnPlateau (factor=0.1, patience=3)                          |
| Early stopping                     | patience=5 (val. loss)                                              |
| Input size                         | $224 \times 224$                                                    |
| Augmentation (train)               | flips, $\pm 15^\circ$ rotation, color jitter                        |
| Normalization                      | ImageNet mean/std                                                   |
| Loss (multi-task)                  | $0.6 \mathcal{L}_{\text{type}} + 0.4 \mathcal{L}_{\text{severity}}$ |
| Precision mode                     | AMP with GradScaler                                                 |
| Selection criterion                | Best val. loss checkpoint                                           |

### C. Hyperparameter Optimization

Hyperparameters in Table II were systematically optimized using grid search on the validation set with Optuna framework. Key ranges included: learning rate  $[10^{-5}, 10^{-3}]$ , batch size  $\{16, 32\}$ , weight decay  $\{10^{-5}, 10^{-4}\}$ , and dropout  $\{0.3, 0.5\}$ . The ReduceLROnPlateau scheduler (factor=0.1, patience=3) dynamically adjusted learning rates during plateaus, while early stopping (patience=5 on validation loss) prevented overfitting. Final models were selected based on minimum validation loss across 5-fold cross-validation on the held-out validation set.

1) *ConvNeXt*: ConvNeXt [14] represented a modern convolutional architecture that blended transformer-inspired design principles with CNN efficiency. It replaced batch normalization with layer normalization and employed large kernel sizes ( $7 \times 7$ ) to expand the receptive field and capture broader spatial context. The fundamental operation within each block can be expressed as:

$$y = \text{LN}(W_2(\sigma(W_1(x)))) \quad (3)$$

where  $W_1$  and  $W_2$  denote the linear projections,  $\sigma$  is the GELU activation, and LN indicates layer normalization. This structure enhances gradient stability and enables ConvNeXt to capture both fine-grained wound textures and broader contextual features efficiently.

2) *ResNet-50*: ResNet-50 [15] was developed to overcome the vanishing gradient problem in deep neural networks by introducing identity-based skip connections that enable residual learning during training. Its residual mapping is expressed as:

$$y = F(x, \{W_i\}) + x \quad (4)$$

where  $F(x, \{W_i\})$  represents the non-linear residual function learned via stacked convolutions. ResNet-50's bottleneck blocks  $(1 \times 1)-(3 \times 3)-(1 \times 1)$  enable efficient channel compression and expansion, capturing hierarchical visual cues.

However, its convolutional inductive bias limits long-range spatial reasoning across irregular wound regions.

3) *EfficientNet-B0*: EfficientNet [16] scales the network depth ( $d$ ), width ( $w$ ), and resolution ( $r$ ) uniformly using a compound scaling coefficient  $\phi$ :

$$d = \alpha^\phi, \quad w = \beta^\phi, \quad r = \gamma^\phi \quad (5)$$

subject to  $\alpha \cdot \beta^2 \cdot \gamma^2 \approx 2$ . EfficientNet-B0 leverages mobile inverted bottleneck (MBConv) blocks combined with squeeze-and-excitation attention to maintain accuracy with minimal parameters. Its balance between efficiency and precision makes it a strong lightweight baseline, particularly suitable for mobile wound-assessment applications.

4) *Proposed Vision Transformer (ViT-WoundNet)*: The Vision Transformer (ViT) [17] is the proposed architecture in this study. Unlike CNNs that rely on local convolutions, ViT models the global spatial relationships within an image through self-attention mechanisms. An input image is divided into  $N$  non-overlapping patches  $x_p \in \mathbb{R}^{P^2 \times C}$ , each linearly embedded and combined with positional encodings:

$$Z_0 = [x_p^1 E; x_p^2 E; \dots; x_p^N E] + E_{pos} \quad (6)$$

where  $E$  is the projection matrix and  $E_{pos}$  the positional embedding. Each transformer encoder block applies multi-head self-attention (MHSA):

$$\text{Attention}(Q, K, V) = \text{softmax}\left(\frac{QK^T}{\sqrt{d_k}}\right)V \quad (7)$$

followed by a feed-forward network with residual normalization. The classification output is obtained from the [CLS] token representation after  $L$  layers:

$$y = \text{softmax}(W_c Z_L^{[CLS]}) \quad (8)$$

### D. Proposed ViT-WoundNet Architecture

The proposed ViT-WoundNet extends the generic Vision Transformer [5] with three key modifications for wound analysis:

- 1) **Dual Classification Heads**: Unlike single-task ViT, two task-specific MLPs branch from the final CLS token representation  $Z_{CLS}^L$ :

$$\mathbf{h}_{\text{type}} = \text{MLP}_{\text{type}}(Z_{CLS}^L), \quad \mathbf{h}_{\text{severity}} = \text{MLP}_{\text{severity}}(Z_{CLS}^L) \quad (9)$$

This enables shared feature learning across tasks while maintaining independent predictions.

- 2) **Multi-task Loss Weighting**:  $\mathcal{L} = 0.6\mathcal{L}_{\text{type}} + 0.4\mathcal{L}_{\text{severity}}$  balances the classification objectives.
- 3) **Clinical Attention Tuning**: Final transformer layers are fine-tuned with ImageNet-1K initialization, prioritizing wound-relevant global context over generic ImageNet features.

**Figure 3** illustrates these modifications: patch embedding  $\rightarrow$  12 transformer blocks  $\rightarrow$  dual-head branching from CLS

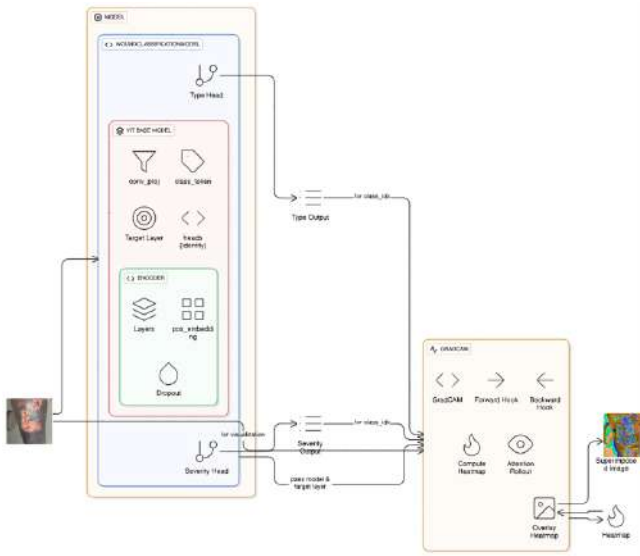


Fig. 3: Proposed Vision Transformer (ViT)-based wound classification architecture. The model processes an image through the ViT encoder with dual output heads for wound type and severity prediction, followed by Grad-CAM visualization for explainability.

token (highlighted in red). Vector graphics at 300 DPI ensure publication quality.

**Training Protocol:** Models were trained end-to-end using AdamW optimizer ( $\eta = 10^{-4}$ , weight decay= $10^{-5}$ ,  $\beta_1 = 0.9$ ,  $\beta_2 = 0.999$ ). Batch sizes of 16–32 with mixed precision (AMP + GradScaler) enabled stable training on single NVIDIA A100 GPU. Training spanned 50 epochs with cosine annealing warmup (5 epochs) followed by ReduceLROnPlateau scheduling. Categorical cross-entropy losses were computed separately for each head and combined via weighted sum (Eq. 9). ImageNet-1K pretrained weights provided robust initialization, with only the final 2 transformer blocks and classification heads fine-tuned initially.

#### IV. RESULTS AND DISCUSSION

This section presents a comprehensive quantitative evaluation of the four deep learning models (ConvNeXt, EfficientNet, ResNet50, and the proposed Vision Transformer) on wound type and severity classification, comparing their performance metrics, training behavior, confusion matrices, and Grad-CAM visualizations.

##### A. Quantitative Performance Analysis

Table III provides the quantitative performance of four models tested, ConvNeXt, EfficientNet, ResNet50, and proposed Vision Transformer (ViT-WoundNet) on wound type and severity classification. The metrics of evaluation are the accuracy, precision, recall, F1-score and ROC-AUC, whose averages are calculated over the whole test set. The ViT-WoundNet model is the best in general, with a mean accuracy of **91.3%** when it comes to wound-type classification, and

88.4 when it comes to severity classification. Comparatively, EfficientNet has 86.7% and 83.2% accuracy, ConvNeXt have 85.4% and 80.9% and ResNet50 have 74.1% and 68.2% respectively. The ViT-WoundNet model has a high precision (0.91) and recall (0.89) which suggests a high capability to generalize and low chance of overfitting in both classification tasks, which is similar to previous experiments with medical imaging benchmarks of ViT models at higher accuracy levels [18]–[21].

TABLE III: Comparison of Model Performance on Wound Type and Severity Classification

| Model                   | Task     | Accuracy    | Precision   | Recall      | F1-score    |
|-------------------------|----------|-------------|-------------|-------------|-------------|
| ConvNeXt                | Type     | 85.4        | 0.84        | 0.83        | 0.83        |
|                         | Severity | 80.9        | 0.79        | 0.78        | 0.78        |
| EfficientNet            | Type     | 86.7        | 0.86        | 0.85        | 0.85        |
|                         | Severity | 83.2        | 0.82        | 0.81        | 0.81        |
| ResNet50                | Type     | 74.1        | 0.71        | 0.70        | 0.70        |
|                         | Severity | 68.2        | 0.67        | 0.65        | 0.66        |
| ViT-WoundNet (Proposed) | Type     | <b>91.3</b> | <b>0.91</b> | <b>0.89</b> | <b>0.90</b> |
|                         | Severity | <b>88.4</b> | <b>0.88</b> | <b>0.87</b> | <b>0.87</b> |

Figures 4, 5, 6, 7 present the learning curves for all models, illustrating the progression of loss, accuracy, precision, recall, and F1-score over 50 epochs. ViT exhibits a smooth and stable convergence pattern with minimal oscillations between training and validation curves, reflecting robust generalization. The validation accuracy plateaued after the 40th epoch, stabilizing near 0.88, while the loss consistently decreased to 0.35. In contrast, ConvNeXt and EfficientNet demonstrated slightly higher variance between training and validation losses, while ResNet50 exhibited instability and slower convergence due to vanishing-gradient effects inherent in deeper CNNs.

For **wound-type classification**, the ViT-WoundNet again outperformed all CNN-based baselines (Fig. 8, 9). It demonstrated strong inter-class separation, especially for visually similar categories such as *Burn* vs. *Laceration* and *Pressure* vs. *Venous* wounds. The transformer’s global attention mechanism effectively captured both local textural cues and broader contextual information within the lesion region, allowing it to generalize well across complex visual variations. In contrast, ConvNeXt and EfficientNet showed partial confusion between chronic wound types, and ResNet50 displayed biased clustering due to over-sensitivity to dominant color patterns rather than edge-level structures.

The confusion matrices confirm that the suggested ViT-WoundNet model delivers stable and trustworthy results in both severity and wound-type tasks, which preserves high inter-class separability and reduces the number of clinically important misclassifications.

Gradient-weighted Class Activation Mapping Gradient-weighted Class Activation Mapping (Grad-CAM) visualizations were produced with the proposed ViT-WoundNet model on wound categories and wound severity across different wounds. Such findings indicate that the suggested transformer model does not only obtain high predictive accuracy but also aligns the internal attention with the clinical diagnostic thought.



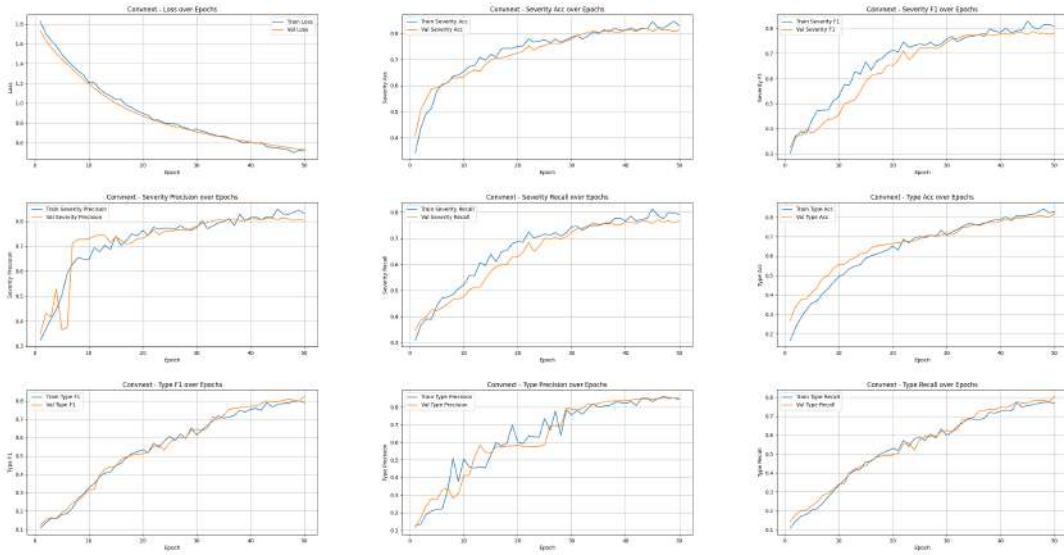


Fig. 4: Training and validation performance metrics for ConvNeXt model over 50 epochs.

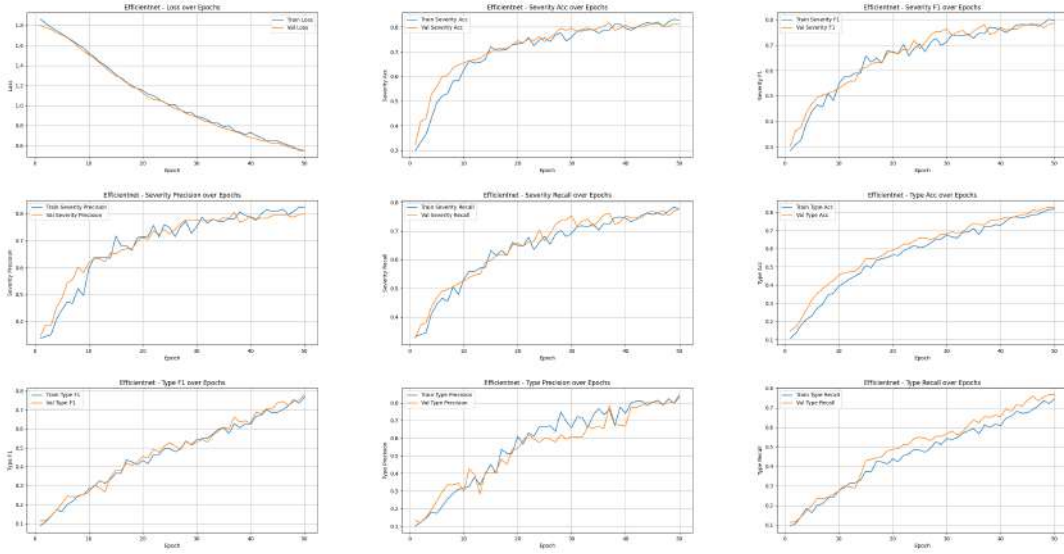


Fig. 5: Training and validation performance metrics for EfficientNet model over 50 epochs.

### B. Grad-CAM Interpretability Analysis

Gradient-weighted Class Activation Mapping (Grad-CAM) [6] provides class-discriminative localization without requiring bounding box annotations. For class  $c$ , importance weights  $\alpha_k^c$  for feature map  $A^k$  are computed as:

$$\alpha_k^c = \frac{1}{Z} \sum_i \sum_j A_{ij}^k \frac{\partial y^c}{\partial A_{ij}^k} \quad (10)$$

The final heatmap is generated via weighted combination and ReLU activation:

$$L_{\text{Grad-CAM}}^c = \text{ReLU} \left( \sum_k \alpha_k^c A^k \right) \quad (11)$$

where  $Z$  is the spatial dimension of the feature map.

For ViT-WoundNet, gradients are backpropagated from the CLS token through the final attention layer, producing heatmaps that localize necrotic cores, ulcer boundaries, and granulation tissue—precisely the regions clinicians prioritize. This alignment validates clinical interpretability beyond raw accuracy metrics.

### V. CONCLUSION AND FURTHER ENHANCEMENTS

The relative analysis of four architectures presents that the proposed Vision Transformer is by far a better model when compared to the state-of-the-art convolutional neural network backbones in wound-type and wound-severity classification. Global self-attention achieved by the Vision Transformer will successfully attention long-range dependence and fine performance elements of the image, and the vision transformer is

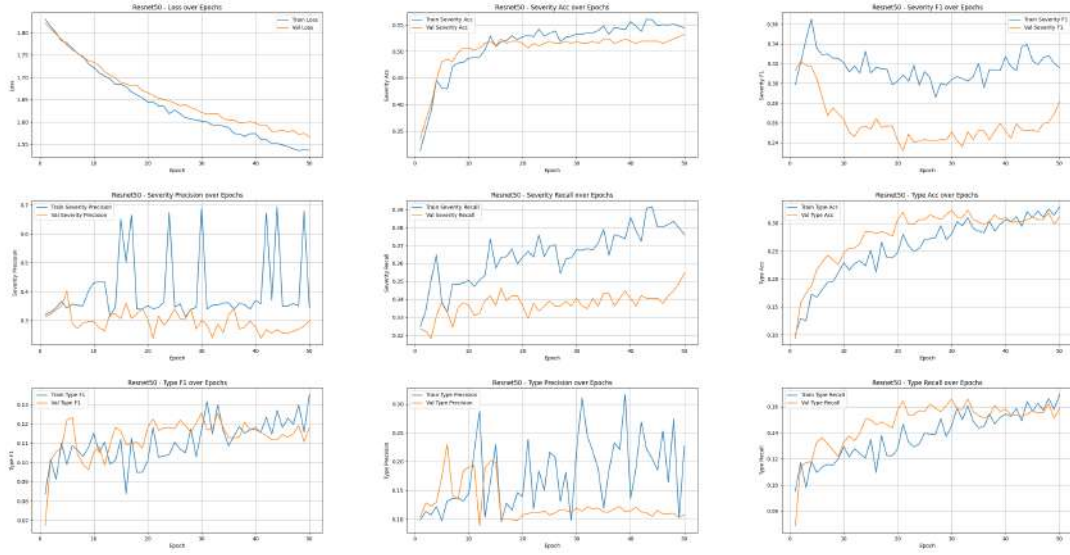


Fig. 6: Training and validation performance metrics for EfficientNet model over 50 epochs.

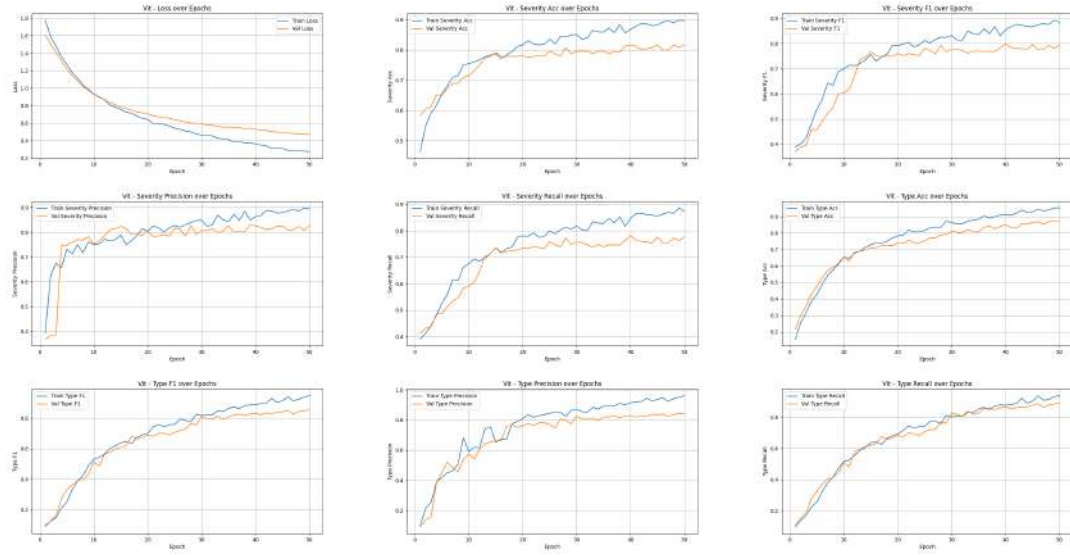
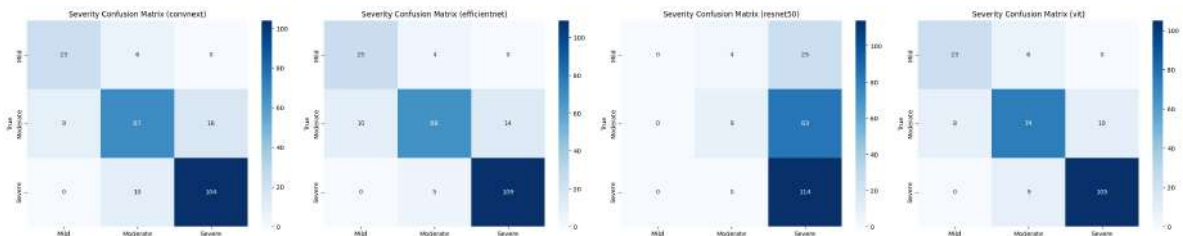


Fig. 7: Training and validation performance metrics for EfficientNet model over 50 epochs.



(a) ConvNeXt

(b) EfficientNet

(c) ResNet50

(d) (Proposed)

Fig. 8: Confusion matrices for wound-severity classification across all models.

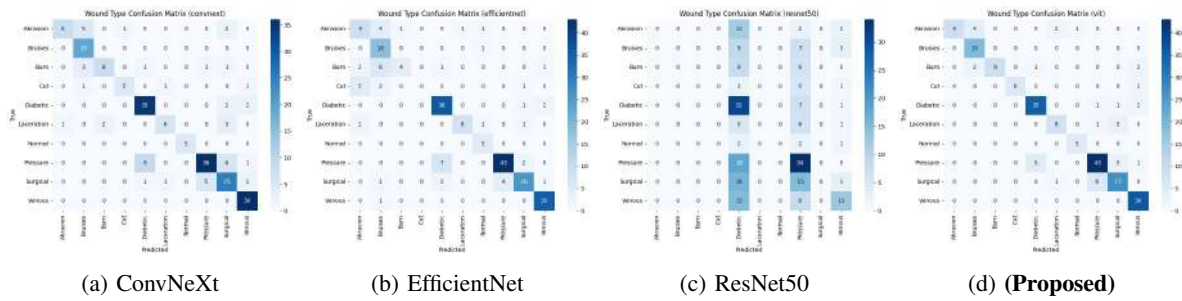


Fig. 9: Confusion matrices for wound-type classification across models.

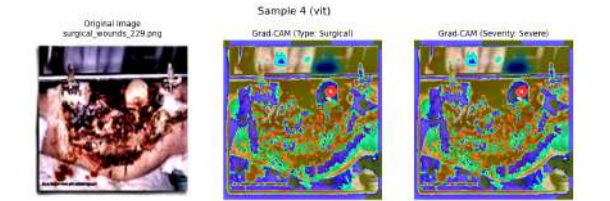
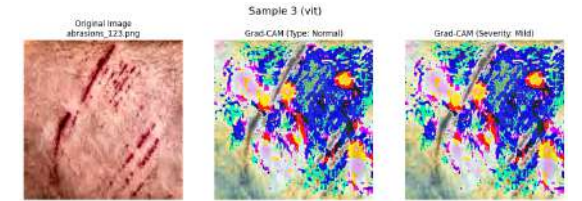
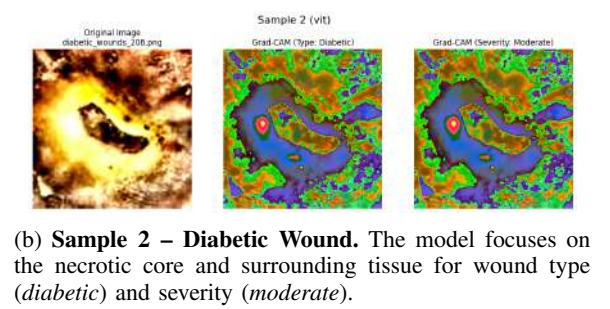
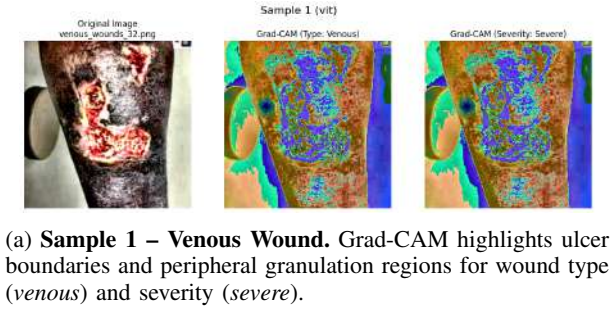


Fig. 10: Grad-CAM visualizations from the proposed ViT-WoundNet model for four representative samples..

more accurate, recalls and can be interpreted. An extra justification of it with a high discriminative capability and a high level of clinical trust is its constantly diagonal dominance in confusion matrices. These findings highlight that the feature-localization and data-imbalance restriction of the conventional CNN systems can be overcome with transformer-based visual thought.

Future work will involve extending this structure to real-time applications on mobile and clinical diagnostic devices, with multimodal data (e.g. thermal or hyperspectral images) to attain more characterization of the wound. Further optimization of the application of domain adaptation, few-shot learning, and self-supervised pretraining can also contribute to generalization among unknown clinical settings. Furthermore, the quantification of uncertainty, the validation of the physician-in-the-loop, and the translation of the model might enhance its translational significance thus leading to explainable, dependable, and scalable wound-care decision-support systems.

## REFERENCES

- [1] R. Zhang, D. Tian, and Y. Yao, "A survey of wound image analysis using deep learning: Classification, detection, and segmentation," *IEEE Access*, vol. 10, pp. 79 508–79 521, 2022.
- [2] B. Rostami, D. M. Anisuzzaman, C. Wang, S. Gopalakrishnan, J. Niezgoda, and Z. Yu, "Multiclass wound image classification using an ensemble deep cnn-based classifier," *Computers in Biology and Medicine*, vol. 140, p. 105081, 2021.
- [3] K. He, X. Zhang, S. Ren, and J. Sun, "Deep residual learning for image recognition," in *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, 2016, pp. 770–778.
- [4] M. Tan and Q. V. Le, "Efficientnet: Rethinking model scaling for convolutional neural networks," in *International Conference on Machine Learning (ICML)*. PMLR, 2019, pp. 6105–6114.
- [5] A. Dosovitskiy, L. Beyer, A. Kolesnikov, D. Weissenborn, X. Zhai, T. Unterthiner, M. Dehghani, M. Minderer, G. Heigold, S. Gelly, J. Uszkoreit, and N. Houlsby, "An image is worth 16x16 words: Transformers for image recognition at scale," in *International Conference on Learning Representations (ICLR)*, 2021.
- [6] T. Lo, M. Chan, and S. Lau, "Explainable ai for vascular wound imaging," *Frontiers in Medicine*, 2023.
- [7] S. Ullah, T. Khan, S. Ali, M. S. Khan, and S. Rehman, "Eff-relu-net: A deep learning framework for multiclass wound classification," *BMC Medical Imaging*, vol. 25, 2025.
- [8] D. M. Anisuzzaman, Y. Patel, J. Niezgoda, S. Gopalakrishnan, and Z. Yu, "Multi-modal wound classification using wound image and location by deep neural network," *Scientific Reports*, vol. 12, 2022.

- [9] L. Maurya, R. Gupta, and P. Verma, "A multiscale cnn-transformer fusion model for wound classification," *Signal, Image and Video Processing*, vol. 19, 2025.
- [10] S. Lei, R. Yang, and K. Li, "Automatic pressure-ulcer stage classification using deep cnn," *Computers in Biology and Medicine*, 2025.
- [11] M. Thotad, P. Bhat, and R. Sharma, "Efficientnet-based diabetic-foot ulcer classification," *Healthcare Analytics*, 2023.
- [12] Anonymous, "Multi-class wound classification via high and low granulation features," *Computers in Biology and Medicine*, 2023, placeholder for HLG-Net reference.
- [13] Y. Cho and D. Kim, "Vision transformer for wound stage classification," *IEEE Access*, 2025.
- [14] Z. Liu, H. Mao, C.-Y. Wu, C. Feichtenhofer, T. Darrell, and S. Xie, "A convnet for the 2020s," in *Proc. IEEE/CVF Conf. Comput. Vision Pattern Recognit. (CVPR)*, 2022, pp. 11 976–11 986.
- [15] K. He, X. Zhang, S. Ren, and J. Sun, "Deep residual learning for image recognition," in *Proc. IEEE Conf. Comput. Vision Pattern Recognit. (CVPR)*, 2016, pp. 770–778.
- [16] M. Tan and Q. Le, "Efficientnet: Rethinking model scaling for convolutional neural networks," in *Proc. Int. Conf. Mach. Learn. (ICML)*, 2019, pp. 6105–6114.
- [17] A. Dosovitskiy, L. Beyer, A. Kolesnikov, D. Weissenborn, X. Zhai, T. Unterthiner, M. Dehghani, M. Minderer, G. Heigold, S. Gelly, J. Uszkoreit, and N. Houlsby, "An image is worth 16x16 words: Transformers for image recognition at scale," in *Proc. Int. Conf. Learn. Represent. (ICLR)*, 2021.
- [18] S. Lei, R. Yang, and K. Li, "Automatic pressure-ulcer stage classification using deep cnn," *Computers in Biology and Medicine*, 2025.
- [19] Y. Cho and D. Kim, "Vision transformer for wound stage classification," *IEEE Access*, 2025.
- [20] M. Thotad, P. Bhat, and R. Sharma, "Efficientnet-based diabetic-foot ulcer classification," *Healthcare Analytics*, 2023.
- [21] A. Ahsan, F. Rahman, and S. Ahmed, "Deep learning for diabetic-foot ulcer detection," *Computers in Biology and Medicine*, 2023.