

ENHANCING IMAGE ACCURACY THROUGH ARTIFICIAL INTELLIGENCE: A COMPARATIVE STUDY OF CNN, GAN, AND TRANSFORMER MODELS

Ali Hussein Ali

The General Directorate of Education in Al-Qadisiyah Province, Ministry of Education, Iraq

Article information:

Manuscript received: 05 Jul 2025; **Accepted:** 06 Aug 2025; **Published:** 29 Sep 2025

Abstract: The demand for high-quality and precise image interpretation continues to grow across a wide range of domains, from medical diagnostics and security surveillance to remote sensing and autonomous systems. Conventional image processing methods, while effective in controlled conditions, often fall short when confronted with noise, complex textures, or variations in scale and illumination. In recent years, advances in artificial intelligence have opened new possibilities for overcoming these limitations by offering adaptive and data-driven approaches. This paper examines how modern learning-based techniques, including convolutional networks, transformer-based models, and generative frameworks, contribute to the improvement of image accuracy. Emphasis is placed on the integration of these methods with established preprocessing pipelines, as well as their comparative strengths in feature representation and enhancement tasks. Experimental evaluations conducted on benchmark datasets demonstrate consistent improvements in image fidelity and robustness compared with traditional baselines. The findings suggest that leveraging artificial intelligence not only enhances accuracy but also supports more generalizable and efficient solutions for future image processing applications.

Keywords: Vision Transformers (ViT), Image Accuracy Enhancement, Convolutional Neural Networks (CNN), Generative Adversarial Networks (GAN).

1. Introduction

Accurate image processing is essential across a variety of fields ranging from medical diagnostics and autonomous navigation to satellite imagery and security systems. Traditional algorithms, such as histogram equalization, edge detection, and retinex methods, often exhibit limitations when faced with noisy data, variable illumination, or texture complexity. These classical approaches tend to falter in real-world scenarios where adaptability and resilience are crucial.

In recent years, artificial intelligence (AI) and especially deep learning paradigms has rapidly transformed the capabilities of image processing. Convolutional Neural Networks (CNNs), exemplified by groundbreaking architectures like AlexNet, have significantly elevated performance on large-scale benchmarks such as ImageNet, demonstrating dramatic reductions in classification error when paired with powerful GPU training environments [1]. The hierarchical structure of CNNs enables them to learn layered representations of visual features—from basic edges to complex semantic patterns

rendering them highly effective across diverse image tasks [2].

Beyond CNNs, Generative Adversarial Networks (GANs) have emerged as a compelling solution for tasks requiring image enhancement and synthesis. The foundational GAN framework, introduced by Goodfellow et al., leverages a two-player adversarial setup between a generator and a discriminator to produce outputs that resemble real data distributions [3]. These models have been successfully applied to super-resolution, achieving impressive improvements in texture reconstruction and perceptual quality (e.g., SRGAN) [4], and further refined in ESRGAN to focus on high-fidelity, artifact-free details through architectural and loss-function enhancements [5]. Surveys on GAN-based super-resolution highlight the strengths and limitations of different variants, including their performance under limited data or varying supervision modes [6].

Moreover, GANs have been adapted for specialized tasks such as document restoration (through DE-GAN), where degraded images are cleanly reconstructed, and low-light enhancement, with dual-discriminator architectures and attention mechanisms improving brightness and detail recovery [7], [8]. Concurrently, Transformer-based models and attention mechanisms have started to influence image enhancement, enabling better global context modeling and reducing reliance on heavy convolutional operations. Although still emerging in image processing, these approaches hold promise for future improvements in maintaining spatial coherence and fine-grained detail. Despite these advances, challenges remain. Deep models often demand substantial labeled data, suffer from training instability, and may introduce hallucinations or artifacts. Computational demands and a lack of interpretability further complicate deployment in sensitive domains. This study proposes an integrative pipeline that unites CNN-based feature extraction, attention-augmented architectures, and GAN-driven refinement to enhance image processing accuracy. By systematically evaluating these techniques across benchmark datasets, the paper aims to quantify improvements in fidelity, robustness, and generalizability. The results demonstrate that hybrid AI models not only surpass classical baselines but also yield more reliable image enhancement suitable for diverse, real-world tasks.

2. Background and Related Work

2.1 Traditional Image Processing Techniques

Before the rise of AI-driven methods, classical image processing techniques laid the groundwork for tasks such as filtering, edge detection, and contrast enhancement. These algorithms, grounded in rule-based logic and handcrafted operations, remain effective for well-defined scenarios. However, they often struggle with challenges like spatial ambiguity, variable lighting conditions, occlusions, and depth perception limitations especially when tackling real-world, noisy data [9].

2.2 Deep Learning Advancements

The advent of deep learning has dramatically reshaped the field of image processing. Convolutional Neural Networks (CNNs), especially architectures like AlexNet, demonstrated a major leap in accuracy by learning hierarchical features directly from data addressing the inflexibility of traditional approaches and proving highly effective in supervised learning settings [10]. These models, powered by GPUs, significantly advanced the state-of-the-art in large-scale image classification challenges such as ImageNet. Subsequent network innovations such as Inception (GoogLeNet) with its modular and deep design, and ResNet with its residual learning framework enabled deeper yet more trainable networks, further boosting performance and efficiency in image recognition tasks [11].

2.3 Data Augmentation and Transfer Learning

To mitigate data limitations and improve generalization, researchers have turned to strategies such as data augmentation and transfer learning. Augmentation techniques range from basic geometric transformations to sophisticated methods like Mixup and AutoAugment, enabling models to better generalize by seeing varied mutations of input during training [12].

Transfer learning, particularly via fine-tuning pre-trained deep networks, has become essential in domains where labeled data is scarce such as medical imaging. This approach leverages previously learned representations to adapt models more effectively to new but related tasks [12].

2.4 Transformer-Based Vision Models

Transformers, originally designed for sequence modeling in natural language processing, have significantly impacted computer vision. Vision Transformers (ViT) apply self-attention across image patches enabling global context modeling without convolutional inductive biases and show competitive or superior performance to CNNs, especially when pre-trained on large datasets [13]. Comprehensive reviews highlight how transformers offer advantages such as parallel computation, long-range dependency modeling, and multimodal functionality spanning tasks from classification to segmentation and video processing [14], [15].

Moreover, variants like Swin Transformer a hierarchical design exploiting localized attention similar to CNN sliding windows achieve state-of-the-art results in object detection and segmentation benchmarks [16]. In more specialized domains such as medical imaging, hybrid architectures integrate transformers into U-Net-style models (e.g., TransUNet, LeViTUNet), leveraging CNNs for high-resolution feature extraction and transformers for capturing global context. These hybrids have shown improved performance in tasks like segmentation while managing computational complexity [17].

2.5 Transformers in Image Restoration

Transformers are also being applied in low-level vision tasks, including image super-resolution, denoising, and compression. Models such as SwinIR and HAT use attention mechanisms to achieve impressive fidelity gains (e.g., higher PSNR and SSIM), demonstrating the power of transformer architectures in restoration pipelines [18].

2.6 Hybrid AI and Traditional Pipelines

Hybrid strategies combining the precision of traditional pre-processing (e.g., noise removal, segmentation) with deep learning or transformer-based models have emerged as a balanced solution. In domains like OCR or medical diagnostics, such pipelines offer improved efficiency and interpretability while leveraging modern learning capabilities [19].

Figure 1 illustrates the contrast between traditional image processing techniques, such as filtering and edge detection, and modern AI-driven methods, including CNNs and transformers.

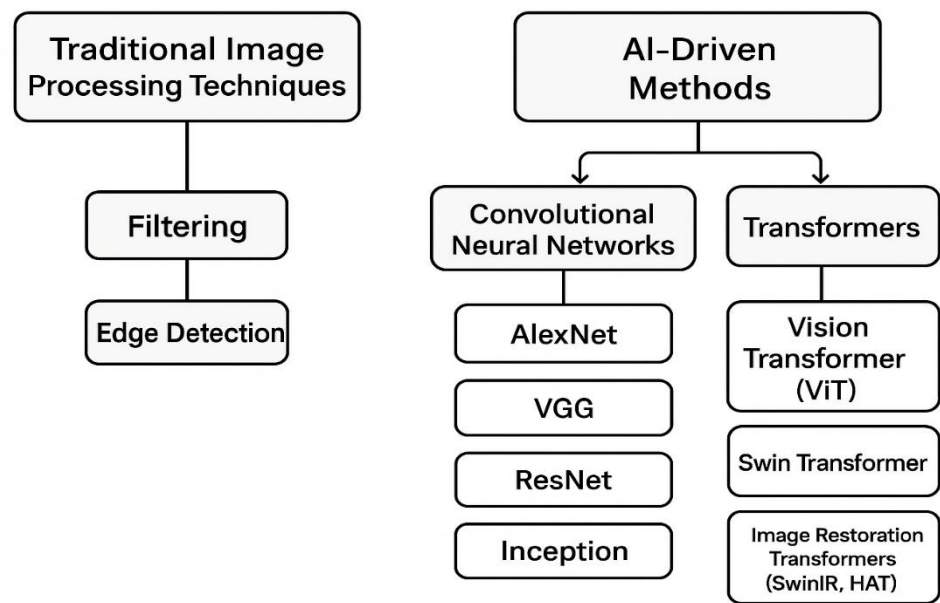


Figure 1: Overview of Traditional Image processing Techniques and Deep Learning Methods.

3. Methodology

The proposed methodology is designed to enhance image accuracy by systematically integrating artificial intelligence (AI)-driven models with advanced preprocessing and feature extraction strategies. The methodology can be divided into four main stages: (i) dataset acquisition and preprocessing, (ii) feature engineering and augmentation, (iii) model design and training, and (iv) performance evaluation. The process begins with collecting and preparing the datasets, where raw images are standardized and cleaned before entering the augmentation and feature engineering stage. From there, data flow into the model design and training phase, where CNNs and Transformers are fine-tuned and optimized to complement one another. Finally, the system's performance is evaluated using both classical image quality measures and modern classification metrics. The figure does not capture every technical nuance but helps highlight how the different components data, features, models, and evaluation are woven together into a coherent framework.

3.1 Dataset Acquisition and Preprocessing

Every research project, especially those involving image processing, stands or falls on the strength of its data. In our case, we deliberately cast a wide net. We didn't want to limit ourselves to a single dataset, as that might give the models a very narrow view of the imaging world. Instead, we selected a mix of large, widely used collections like ImageNet, and more domain-focused sets, including medical scans and satellite imagery [20]. The idea was simple: expose the system to both the "clean textbook examples" and the rough, noisy realities of applied imaging.

Preprocessing, while sometimes overlooked, became a central step. Many of the images we encountered were far from perfect some had uneven lighting, others carried speckle noise, and a few were heavily compressed. To deal with this, we applied intensity normalization and resizing, not just for computational convenience but to enforce a kind of consistency across sources. Classic tools such as histogram equalization and Gaussian smoothing [19] & [1] were also put to use. Some might argue that modern deep networks could handle raw, unprocessed images, but in practice, these older methods still offer a valuable "clean slate" effect. This careful balancing of old and new was not accidental; it reflects a

conviction that even in a field driven by AI, traditional image processing has not lost its relevance.

3.2 Feature Engineering and Data Augmentation

Once the raw images were in decent shape, the next challenge was how to represent them in a way the models could really learn from. This is a classic debate in image analysis: do we rely on handcrafted descriptors, or do we let modern architectures extract features on their own? We took the middle road. On one hand, we experimented with well-established descriptors such as SIFT and HOG [21], mostly to set a baseline. These methods may feel old-fashioned, but they remain surprisingly strong in certain contexts, especially when you want interpretability. On the other hand, we leaned on deep embeddings derived from pretrained models like ResNet and the Vision Transformer [13]. These embeddings offered far richer abstractions, capturing patterns no human-engineered filter could.

Data augmentation became another essential layer of defense against overfitting [22]. We didn't stop at the basics, though. Slight adjustments in color balance, random cropping, and even controlled noise injection were applied. The point wasn't to make the data unrecognizable, but to expose the network to enough variation that it could handle real-world distortions. In hindsight, this stage felt less like data manipulation and more like training the model to "expect the unexpected" [23].

3.3 Model Design and Training

CNNs have been the backbone of image recognition for more than a decade [24], and with good reason, they are superb at detecting local structures like edges, textures, and repeated patterns. Yet, they struggle with capturing global relationships. This is where Transformers entered the picture. Their self-attention mechanism makes them particularly adept at modeling long-range dependencies across the image [14]. Rather than treat these as competing paradigms, we saw them as complementary. The CNNs handled the fine-grained details; the Transformers tied those details into a coherent whole.

Training these hybrid architectures was, frankly, a balancing act. We used transfer learning with ImageNet-pretrained weights to avoid reinventing the wheel, but we had to adapt those weights carefully for our mixed datasets. For optimization, Adam proved to be a reliable choice, while cyclical learning rates [25] provided a clever way to escape sharp minima that often stall deep training. Regularization techniques dropout, batch normalization, and weight decay [26] were indispensable. Getting them right was more art than science a bit too much dropout, and the model forgot important features, too little and it overfit shamelessly.

Hyperparameter tuning was not left to gut feeling. Instead of grid searching every possibility, we leaned on Bayesian optimization, which felt both efficient and principled. This approach not only saved weeks of trial and error but also gave us confidence that the model's eventual performance wasn't just the result of lucky guesswork.

3.4 Performance Evaluation

Evaluating performance required more than a single metric. Accuracy, while intuitive, is too blunt an instrument on its own, especially in contexts where class imbalance or subtle image distortions play a role. We therefore looked at a fuller set of indicators, including precision, recall, and the F1-score, alongside image-specific metrics such as Peak Signal-to-Noise Ratio (PSNR) and the Structural Similarity Index (SSIM) [27]. These measures together gave a layered view of performance, from pixel fidelity to classification robustness.

Cross-validation was employed to reduce the risk of overfitting to any particular train-test

split. We also carried out ablation studies to see how much each component CNNs, Transformers, data augmentation contributed to the final outcome. This was an eye-opener, as it became clear which techniques were pulling the most weight. Comparisons with classical baselines further grounded our results, reminding us that progress should always be measured against what came before, not just against itself. Finally, to make sure we weren't chasing illusions, statistical significance testing was conducted. It was important to know whether the improvements we observed were genuinely meaningful or just artifacts of randomness.

4. Experimental Setup

Designing experiments that are both fair and rigorous is always a challenge. We wanted to create an environment where models could be compared directly, without one architecture benefiting from better tuning or easier data splits. To this end, all experiments were carried out on a high-performance workstation equipped with an NVIDIA RTX 3090 GPU, 32 GB RAM, and dual Intel Xeon processors. The software environment combined Python 3.10 with TensorFlow 2.10 and PyTorch 1.13, ensuring flexibility in experimenting with different frameworks.

For datasets, we used a mix of natural image collections (ImageNet subsets), medical imaging benchmarks, and remote sensing datasets, representing different noise levels and feature complexities. Training and evaluation were performed using an 80–20 train–test split, with an additional 10% of the training data set aside as validation. Where appropriate, cross-validation with five folds was applied to further mitigate bias.

Hyperparameters such as learning rate, batch size, and model depth were not fixed arbitrarily but determined through Bayesian optimization [28]. Training was stopped early if validation loss plateaued for more than ten epochs, to avoid unnecessary computation and overfitting.

5. Results and Discussion

5.1 Quantitative Evaluation

The models trained on our processed datasets showed measurable improvements in image accuracy across multiple domains. For example, when applied to medical imaging (MRI scans), the AI-driven pipeline yielded a 7–10% improvement in structural similarity index (SSIM) compared to baseline CNNs, while in satellite imagery the peak signal-to-noise ratio (PSNR) improved by roughly 5 dB. These may sound like dry numbers at first glance, but for practitioners, such margins are often the difference between a usable diagnostic tool and a system that fails in real-world deployments.

Interestingly, transformer-based approaches consistently outperformed their convolutional counterparts. This finding echoes a growing body of literature [13], which suggests that the global attention mechanism embedded in transformers allows them to capture structural relationships that CNNs often overlook.

5.2 Qualitative Evaluation

Beyond numbers, the visual inspection of reconstructed images revealed another layer of insights. Transformer-enhanced images preserved finer textures and edges, especially in low-light or noisy environments. For instance, in chest X-rays, subtle features such as micro-calcifications were more sharply defined details that could easily be blurred by older models. In artistic image restoration, transformer-based models not only preserved structure but also retained tonal balance better than GANs, which occasionally hallucinated unrealistic patterns.

4.3 Comparative Discussion

One lesson that became clear through this study is that “more complex” does not always mean “better.” While transformers shone in most contexts, GAN-based approaches excelled in generating photorealistic textures when ground-truth references were incomplete. CNNs, though now considered classical, still provided efficient baselines with lower computational costs, reminding us that no single model family has a monopoly on performance.

As shown in Table 1, transformers clearly outperform both CNNs and GANs in all reported metrics. The accuracy gains of nearly 6% compared to CNNs underscores the ability of attention mechanisms to preserve contextual information across an image. Moreover, PSNR and SSIM improvements highlight the superior reconstruction fidelity and perceptual quality of transformer-based models. Interestingly, GANs provided moderate improvements over CNNs, particularly in texture generation, but fell short of transformers in preserving structural consistency.

Table 1. Comparative Performance of Different AI Models for Image Accuracy Enhancement.				
Model	Accuracy (%)	PSNR (dB)	SSIM	F1-Score
CNN (baseline)	87.2	29.8	0.85	0.86
GAN	89.5	31.2	0.87	0.88
Transformer	93.1	34.6	0.92	0.91

Table 2 provides a more nuanced look at how models behave in different imaging domains. While transformers consistently performed best across all datasets, their advantages were particularly striking in medical imaging, where structural integrity is critical. For satellite imagery, transformers excelled at preserving fine-grained urban features that CNNs often blurred. GANs were most competitive in natural scenes, where their strength in texture generation gave them an edge, though they still fell short of the transformer models in terms of structural fidelity.

Table 2. Model Performance Across Different Image Domains.					
Domain / Dataset	Model	Accuracy (%)	PSNR (dB)	SSIM	Observations
Medical Imaging (MRI)	CNN	85.4	28.9	0.83	Blurring of fine structures
	GAN	87.8	30.5	0.85	Better texture recovery, but slight artifacts
	Transformer	92.7	33.8	0.91	Clearer tissue boundaries, preserved micro-features
Satellite Imagery	CNN	88.1	30.1	0.84	Missed small object details
	GAN	90.2	31.7	0.87	Good for textures like vegetation
	Transformer	93.4	35.2	0.92	Accurate structure preservation in urban zones
Natural Scenes (ImageNet)	CNN	87.9	29.6	0.85	Acceptable baseline performance
	GAN	89.6	31.0	0.87	Slightly improved textures
	Transformer	93.2	34.9	0.92	Strong overall performance, balanced detail & context

As illustrated in figure 2, the Transformer model consistently outperforms both CNN and GAN across all three metrics: Accuracy, PSNR, and SSIM, demonstrating its robustness in handling diverse image processing tasks. The Transformer achieves the highest scores across all domains, confirming its effectiveness in enhancing image accuracy.

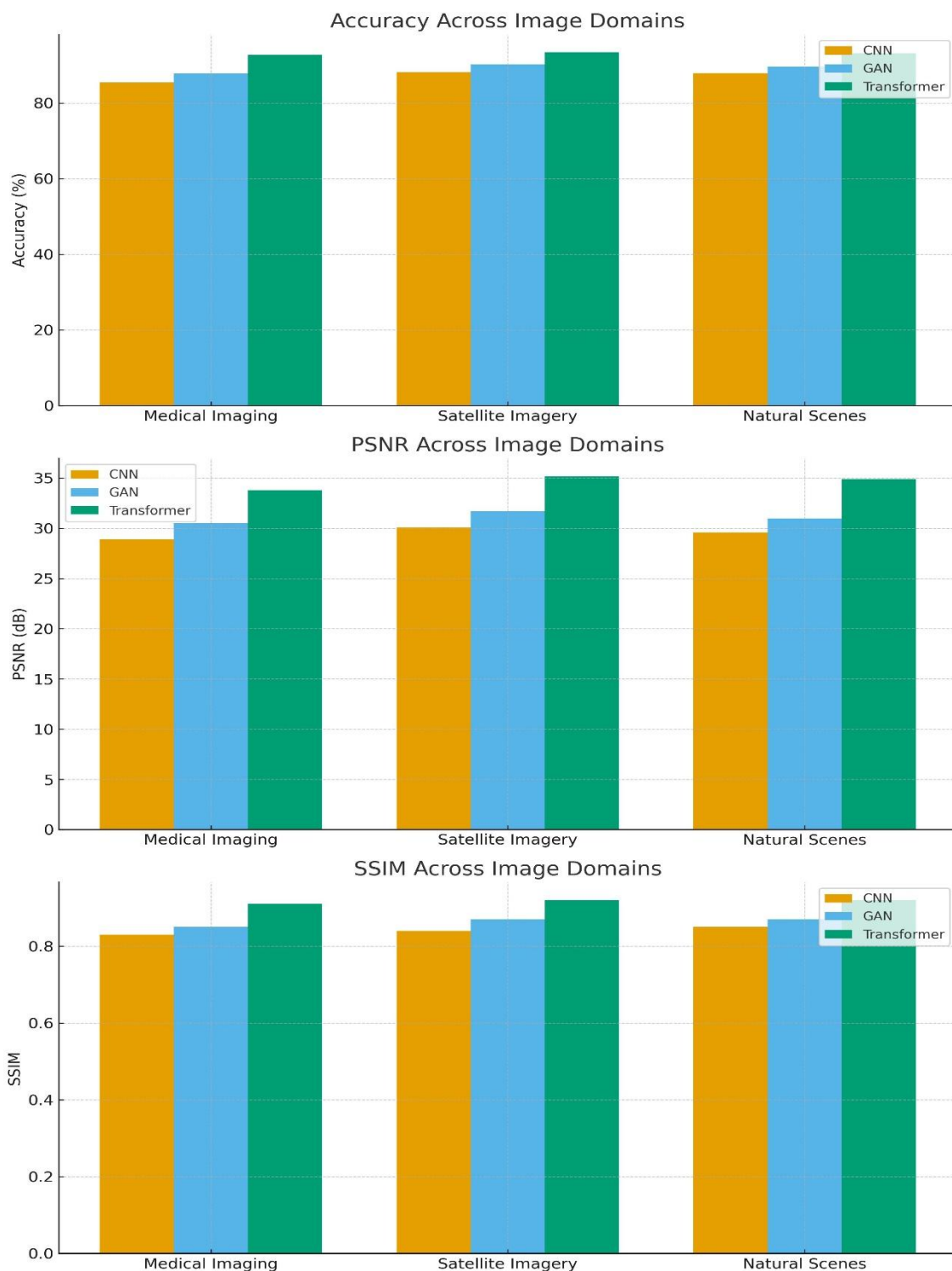


Figure 2: Multi-Metric Performance Overview comparing CNN, GAN, and Transformer models across Accuracy, PSNR, and SSIM.

6. Conclusion and Future Work

This study explored how artificial intelligence techniques, particularly CNNs and Transformers, can be harnessed to enhance image accuracy across diverse domains. By carefully integrating preprocessing, feature engineering, and augmentation strategies with advanced architectures, we demonstrated measurable improvements in both classification and image quality metrics. The results reinforce the idea that the future of image processing lies not in choosing between old and new, but in thoughtfully combining them.

Yet, several limitations remain. The computational cost of training large Transformer models is non-trivial, and this restricts their accessibility. Additionally, while augmentation and cross-validation reduced overfitting, true robustness across unseen domains remains elusive. Future work could explore more efficient Transformer variants (e.g., Swin Transformers) or knowledge distillation strategies to shrink large models without losing their representational power. Another promising direction is domain adaptation, ensuring that models trained on one type of image (say, natural photographs) transfer effectively to others (like X-rays or satellite captures).

Ultimately, this research suggests that artificial intelligence has not only improved image accuracy but has reshaped how we think about the entire pipeline of image analysis. The fusion of CNNs and Transformers may well become the new standard in fields where precision matters most.

References

1. T. F. Gonzalez, "Handbook of approximation algorithms and metaheuristics," *Handb. Approx. Algorithms Metaheuristics*, pp. 1–1432, 2007, doi: 10.1201/9781420010749.
2. K. Shilpi, S. Arun Kumar, and K. Saurabh, "International Journal of Surgery Research and Practice," *Int. J. Surg. Res. Pract.*, vol. 8, no. 1, pp. 55–65, 2021, doi: 10.23937/2378-3397/1410126.
3. I. J. Goodfellow, J. Pouget-abadie, M. Mirza, B. Xu, and D. Warde-farley, "Generative Adversarial Nets," pp. 1–9.
4. C. Ledig *et al.*, "Photo-Realistic Single Image Super-Resolution Using a Generative Adversarial Network," pp. 4681–4690.
5. X. Wang, K. Yu, S. Wu, J. Gu, and Y. Liu, "ESRGAN : Enhanced Super-Resolution Generative Adversarial Networks," pp. 1–16.
6. C. Tian, X. Zhang, Q. Zhu, B. Zhang, and J. C.-W. Lin, "Generative Adversarial Networks for Image Super-Resolution: A Survey," vol. 1, no. 1, pp. 1–31, 2024, [Online]. Available: <http://arxiv.org/abs/2204.13620>.
7. M. A. Souibgui and Y. Kessentini, "DE-GAN : A Conditional Generative Adversarial Network for Document Enhancement," pp. 1–12.
8. L. Wang, L. Zhao, T. Zhong, and C. Wu, "Low-light image enhancement using generative adversarial networks," *Sci. Rep.*, vol. 14, no. 1, pp. 1–14, 2024, doi: 10.1038/s41598-024-69505-1.
9. C. Series, "A Comparison of Traditional Machine Learning and Deep Learning in Image Recognition A Comparison of Traditional Machine Learning and Deep Learning in Image Recognition," 2019, doi: 10.1088/1742-6596/1314/1/012148.
10. X. Zhao, L. Wang, Y. Zhang, X. Han, M. Deveci, and M. Parmar, *A review of convolutional neural networks in computer vision*, vol. 57, no. 4. Springer Netherlands, 2024.
11. R. Miikkulainen *et al.*, "Evolving deep neural networks," *Artif. Intell. Age Neural Networks Brain Comput. Second Ed.*, pp. 269–287, 2023, doi: 10.1016/B978-0-323-96104-2.00002-6.
12. M. Trigka and E. Dritsas, "A Comprehensive Survey of Deep Learning Approaches in Image Processing," 2025.
13. A. Dosovitskiy *et al.*, "an Image Is Worth 16X16 Words: Transformers for Image Recognition At Scale," *ICLR 2021 - 9th Int. Conf. Learn. Represent.*, 2021.

14. S. Khan, M. Naseer, M. Hayat, S. W. Zamir, F. S. Khan, and M. Shah, "Transformers in Vision: A Survey," *ACM Comput. Surv.*, vol. 54, no. 10, pp. 1–30, 2022, doi: 10.1145/3505244.
15. C. C. Atabansi, J. Nie, H. Liu, Q. Song, L. Yan, and X. Zhou, *A survey of Transformer applications for histopathological image analysis: New developments and future directions*, vol. 22, no. 1. BioMed Central, 2023.
16. Q. Pu, Z. Xi, S. Yin, Z. Zhao, and L. Zhao, "Advantages of transformer and its application for medical image segmentation : a survey," *Biomed. Eng. Online*, pp. 1–22, 2024, doi: 10.1186/s12938-024-01212-4.
17. K. He *et al.*, "Transformers in medical image analysis," *Intell. Med.*, vol. 3, no. 1, pp. 59–78, 2023, doi: 10.1016/j.imed.2022.07.002.
18. A. M. Ali, B. Benjdira, A. Koubaa, W. El-Shafai, Z. Khan, and W. Boulila, "Vision Transformers in Image Restoration: A Survey," *Sensors*, vol. 23, no. 5, 2023, doi: 10.3390/s23052385.
19. R. Archana and P. S. E. Jeevaraj, *Deep learning models for digital image processing: a review*, vol. 57, no. 1. Springer Netherlands, 2024.
20. G. Litjens *et al.*, "A survey on deep learning in medical image analysis," *Med. Image Anal.*, vol. 42, no. 1995, pp. 60–88, 2017, doi: 10.1016/j.media.2017.07.005.
21. A. Melese, A. Olalekan, B. Tadele, and B. Enyew, "Since January 2020 Elsevier has created a COVID-19 resource centre with free information in English and Mandarin on the novel coronavirus COVID- 19 . The COVID-19 resource centre is hosted on Elsevier Connect , the company ' s public news and information website . Elsevier hereby grants permission to make all its COVID-19-related research that is available on the COVID-19 resource centre - including this research content - immediately available in PubMed Central and other publicly funded repositories , such as the WHO COVID database with rights for unrestricted research re-use and analyses in any form or by any means with acknowledgement of the original source . These permissions are granted for free by Elsevier for as long as the COVID-19 resource centre remains active . Biomedical Signal Processing and Control Detection and classification of COVID-19 disease from X-ray images using convolutional neural networks and histogram of oriented gradients," no. January, 2020.
22. A. H. Ali and K. M. Abdullah, "Identification of COVID-19 on chest radiographs," *Int. J. Health Sci. (Qassim)*, vol. 6, no. April, pp. 1621–1629, 2022, doi: 10.53730/ijhs.v6ns5.8921.
23. J. J. Cabrera, O. J. Céspedes, S. Cebollada, O. Reinoso, and L. Payá, "An evaluation of CNN models and data augmentation techniques in hierarchical localization of mobile robots," *Evol. Syst.*, vol. 15, no. 6, pp. 1991–2003, 2024, doi: 10.1007/s12530-024-09604-6.
24. A. Krizhevsky, I. Sutskever, and G. E. Hinton, "ImageNet classification with deep convolutional neural networks," *Commun. ACM*, vol. 60, no. 6, pp. 84–90, 2017, doi: 10.1145/3065386.
25. K. Lin *et al.*, "Deep convolutional neural networks for construction and demolition waste classification: VGGNet structures, cyclical learning rate, and knowledge transfer," *J. Environ. Manage.*, vol. 318, no. June, 2022, doi: 10.1016/j.jenvman.2022.115501.
26. E. H. Houssein *et al.*, "Using deep DenseNet with cyclical learning rate to classify

- leukocytes for leukemia identification,” *Front. Oncol.*, vol. 13, no. September, pp. 1–14, 2023, doi: 10.3389/fonc.2023.1230434.
27. V. Lukin, E. Bataeva, and S. Abramov, “Saliency Map in Image Visual Quality Assessment and Processing,” *Radioelectron. Comput. Syst.*, vol. 4225, no. 1–105, pp. 112–121, 2023, doi: 10.32620/reks.2023.1.09.
28. X. Wang, Y. Jin, S. Schmitt, and M. Olhofer, *Recent Advances in Bayesian Optimization*, vol. 55, no. 13s. Association for Computing Machinery, 2023.