

# Analyze the Deep Learning Algorithms for Image Enhancement

Faiza Latif Abbasi<sup>1</sup>, Dilbar Hussain<sup>1\*</sup>, Saira Khurram Arbab<sup>2</sup>

<sup>1</sup>Faculty of Engineering Sciences & Technology (FEST), Iqra University, Main Campus Karachi, 75500, Pakistan  
[faizakhanabbasi@gmail.com](mailto:faizakhanabbasi@gmail.com), [dilbarhussainmalik04@gmail.com](mailto:dilbarhussainmalik04@gmail.com), [sairarbab@iqra.edu.pk](mailto:sairarbab@iqra.edu.pk)

Corresponding Author: [dilbarhussainmalik04@gmail.com](mailto:dilbarhussainmalik04@gmail.com)

**Abstract:** Image reconstruction and image enhancement are basic problems in computer vision, and are of great values for medical image, satellite remote sensing and image restoration in digital photographs. The reconstruction quality is directly related to the diagnostic reliability in medical imaging (CT, MRI, X-ray) and to environmental monitoring or in process decision making in satellite remote sensing, in high stakes applications. Current traditional methods such as interpolation, iterative, and optimization based methods have the disadvantage of being inherently limited with regards to computational efficiency, failure to deal with complex profiles of noise and poor performance on low-resolution input data. This paper introduces an original deep learning architecture that combines a modified 27-layer Convolutional Neural Network (CNN) with the DenseNet121 backbone in a more synergistic way to overcome these shortcomings. The architecture proposed exploits the reuse of dense features and an encoder-decoder structure to capture the local texture features and the global structure, hence, facilitating higher fidelity of the image. The framework is conditioned and tested on the Deep Autoencoder Image Reconstruction benchmark dataset. The results of the experiment show that the proposed method has an accuracy of 94.6%, a precision of 92.3%, a recall of 91%, and an F1-score of 90.3%, which is better than the state-of-the-art conventional and deep learning baselines. The work contributions include the architectural design, the CNN-integration strategy, and a highly rigorous evaluation protocol using multiple metrics. The innovations have great potential in the further evolution of smart imaging systems in a wide range of real-world applications.

**Keywords:** Image Reconstruction, Deep Learning, Deep Autoencoder, Image Enhancement

## 1. INTRODUCTION

Today, image acquisition systems are always subject to a variety of degradation models including sensor noise, motion blur, optical aberrations and hardware resolution limits. These characteristics have a detrimental effect on the information integrity of images taken and are problematic in a number of applications such as medical diagnostics, remote sensing, autonomous navigation and digital archival preservation, where they pose a fundamental challenge to downstream applications [1]. The ability to accurately rebuild or refine damaged images of incomplete, noisy, or low-resolution images is thus a scientific and practical issue of deep interest [2]. In satellite imaging, reconstruction errors can impact environmental monitoring, disaster response and accuracy of surveillance in medical imaging, incorrect reconstruction can create misdiagnosis or the loss of clinically significant information, whereas in satellite imaging, reconstruction errors may impact environmental monitoring, disaster response and the accuracy of surveillance.

The scientific issue that this work focuses on revolves around the fundamental ineffectiveness of classical image reconstruction algorithms, i.e., interpolation-based, iterative and optimization-based algorithms, in the face of complicated, real-world degradation conditions [3]. These traditional methods work based on simplistic assumptions about noise distributions and degradation models, making them inappropriate to heterogeneous corruption profiles observed in practice. What is more, their dependence on hand-written priors and repeatable computational methods penalizes the processing cost, which restricts their use in time-sensitive operation scenarios.

The impetus behind the current study is based on the transformative power of deep learning, and more specifically, Convolutional Neural Networks (CNNs), to learn hierarchical features directly out of data and avoid the reliance on feature engineering on domain-specific features. DenseNet121 specifically has proven to be extremely parameter-efficient due to its dense connectivity structure, which has allowed feature reuse between network layers and greatly reduced the vanishing gradient problem, which afflicts traditional deep architectures. The multi-scale spatial dependencies important to high-fidelity image reconstruction are embodied by the combination of a custom 27-layer CNN and the DenseNet121 backbone, as a principled architectural choice.

The key contributions of the current work are as follows:

- A new hybrid deep learning architecture is suggested, combining a 27-layer CNN with a DenseNet121 backbone into an encoder-decoder framework that is explicitly a reconstruction and enhancement of images.

- The compound loss is designed, which is the combination of Mean Squared Error (MSE), perceptual loss and Total Variation (TV) regularization to optimize pixel-wise fidelity, perceptual similarity, and smoothness jointly.
- The Deep Autoencoder Image Reconstruction benchmark dataset is tested thoroughly using the four widely used accuracy, precision, recall and F1-score metrics of accuracy, and it is found that the proposed methods can be shown to achieve a definite improvement over standard methods.
- The proposed framework is evaluated for its applicability to different areas of imaging including medical imaging (CT, MRI, X-ray), satellite imaging, and photographic restoration.

Photographs that capture historically important events and individual memories are vulnerable to physical degradation when stashed over a long time period, especially when stored in unfavorable conditions. Figure 1 demonstrates the schematic results of the deep learning-based restoration used in the old and degraded photographic material, highlighting the practical potential of the suggested methodology. Moreover, interpolation algorithms tend to give overly smooth results without the high-frequency content of the original, whereas optimization-based methods are computationally expensive and sensitive to parameter tuning, and thus are not suitable in real-time and large-scale applications.



Fig. 1. Illustrative results of deep learning-based old photograph restoration.

## 2. LITERATURE REVIEW

Image reconstruction has been experiencing a paradigm shift in the last decade, leading to the growth of deep learning algorithms and the existence of large-scale annotated datasets. The subsequent subsections are a systematic discussion of the most relevant developments in imaging modalities and architectural paradigms.

### CNN-based Reconstruction in Medical Radiography

CNNs have become the leading architectural model of reconstruction tasks in various medical imaging modalities. Deep learning has been integrated in Positron Emission Tomography (PET) as either purely data-driven mapping systems or regularization penalties in iterative reconstruction systems. Such techniques enable one-to-one conversion of raw measured signals to high-quality reconstructed images that have less noise and increased spatial resolution [4]. It is important to note that CNN-based architectures have shown interesting results in low-dose PET imaging, where the quality of the image is diagnostically suitable, and the radiation dose is reduced significantly [5]. Figure 2 shows how real images converge to the high-perceptual-fidelity outputs of StyleGAN synthesis.

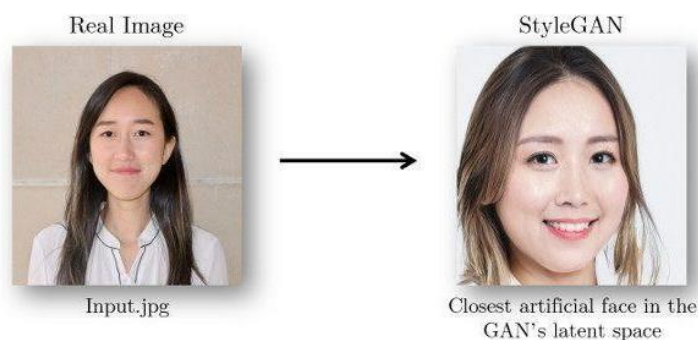


Fig. 2. Real image convergence to StyleGAN Synthesis Image with high quality.

CNNs have found application in the field of super-resolution fluorescence microscopy to speed up the reconstruction of Stochastic Optical Reconstruction Microscopy (STORM) images in order to image live cells. A reconstruction algorithm based on deep CNN and specifically designed to be used in live cell STORM imaging can be used to analyze dynamically changing cellular processes in real-time, with results demonstrating faster reconstructions and better image quality [6]. These properties of CNNs, especially

their ability to obtain complex spatial patterns with high-dimensional datasets, render them especially useful in reconstructing images with complex fine-structure.

Computer Tomography (CT) reconstruction has also been improved by deep learning interventions, with CNN-based interventions showing the capability to generate high-quality 3D volumetric reconstructions at lower radiation doses [7][8]. The use of 3D CNNs in the model of volumetric CT has provided anatomically realistic modeling with higher diagnostic value [9]. Similar progress has been reported in Magnetic Resonance Imaging (MRI), where prior-based reconstruction algorithms based on plug-in deep CNN priors trained on high-quality images can be used to achieve effective reconstruction of under-sampled k-space data, leading to reduced artifact levels and higher image quality [10][11].

### Transfer Learning and Domain Adaptation

Transfer learning (TL) is one such promising technique in medical imaging, where the size of an annotated dataset is frequently costly, unfeasible, and often difficult to access because of privacy concerns and the need to have specialized clinical knowledge [12][13]. TL allows the creation of special-purpose models with significantly lower annotated training data demands by pre-training reconstruction networks on large-scale natural image datasets (like ImageNet) [14][15]. Moreover, TL enables generalization of models across clinical environments, improving the applicability of CNN-based reconstruction systems in practice [16].

### Generative Adversarial Networks Image Reconstruction

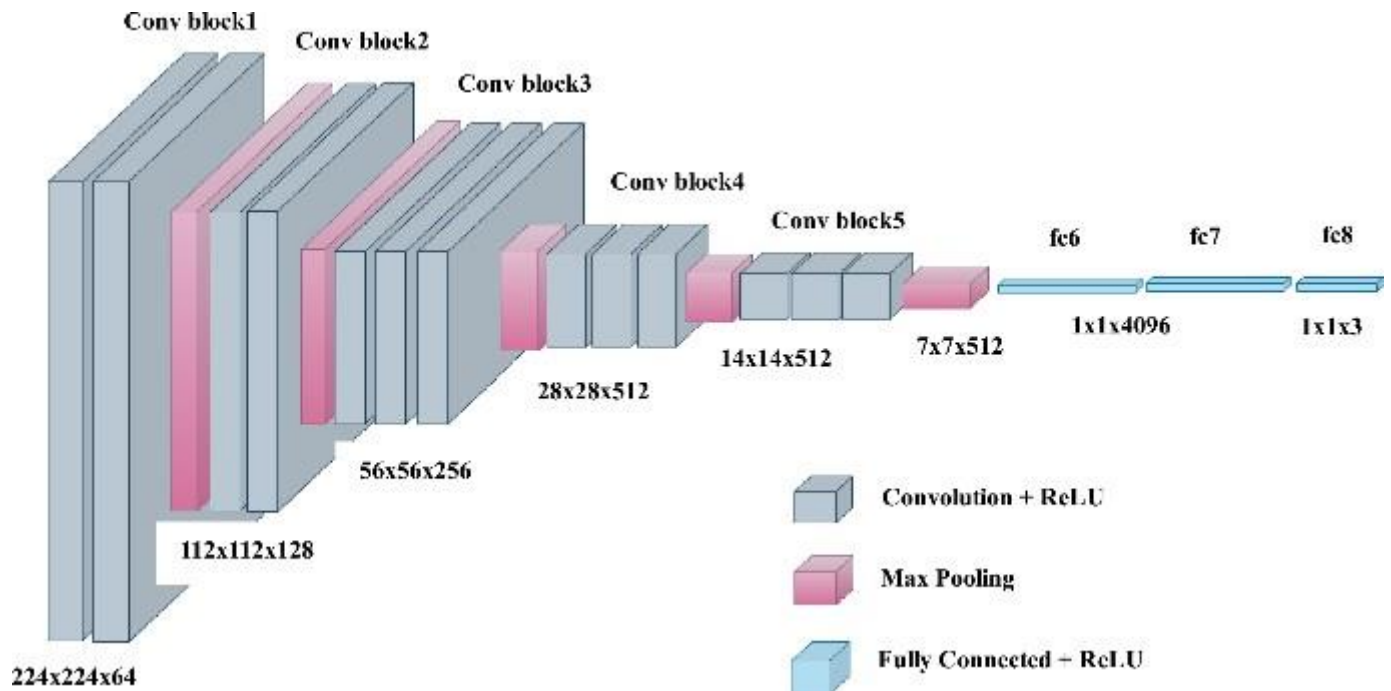
Generative Adversarial Networks (GANs) have already shown a significant potential in image reconstruction, especially in the area of MRI and cross-domain image translation. Image synthesis systems based on GANs can generate high-quality images with under-sampled measurements, with shorter MRI acquisition time without a significant loss in diagnostic quality [17]. Also, GAN models are useful in noise and artifact suppression of low-quality reconstructed images [18]. The weaknesses of GANs compared to CNN-based approaches indicate that a combination of the two can provide additional gains in reconstruction fidelity [19].

### Image Reconstruction Using Vision Transformers

Vision Transformers (ViTs) are a new architectural paradigm in image reconstruction, characterized by self-attention mechanisms to learn long-range spatial dependencies that are beyond the receptive field constraints of traditional convolutional operators [20]. The ViTs were first introduced in 2020 and be especially useful in restoration-related tasks that demand sensitivity to the context of the whole image, such as restoring large areas that have been corrupted. Even though the use of ViTs in medical image reconstruction is still in its infancy, initial results indicate that they can compete with CNN baselines, and incorporation into reconstruction pipelines is a direction to pursue in the future. The GANs combined with Cross-ViT architectures have continued to push the limits of image restoration [21].

## 3. METHODOLOGY

The presented methodology solves the image reconstruction issue by using a systematic pipeline that includes the image dataset preparation, multi-stage preprocessing, architecture design, loss function formulation, and training protocol design. The overall aim is to allow faithful reconstruction of high-quality images with degraded inputs synergistically using CNN-based feature extraction



and DenseNet121-based dense representation learning. The whole workflow of the proposed system is presented in Figure 3.

Fig. 3. Complete workflow of the proposed deep learning image reconstruction system.

## Dataset Description

The proposed model is trained and validated using the Deep Autoencoder Image Reconstruction benchmark dataset and tested on the same dataset. The image data set consists of clean and synthetic degraded image samples that are paired up according to a wide variety of scene categories, image resolutions and image degradations. The dataset partition follows the ratio of 80:10:10 (2/3:1/3:1/3) for training, validation, and testing subsets, respectively, which ensures the statistically robust performance estimation. All images are being captured from public image repositories and then processed to be standardized before being ingested into the model.

## Preprocessing Pipeline

A systematic preprocessing protocol is applied to ensure dataset suitability for deep model training:

- **Resizing:** All images are uniformly resized to  $224 \times 224$  pixels to conform to the DenseNet121 input specification, ensuring architectural compatibility and computational consistency.
- **Normalization:** Pixel intensity values are linearly scaled to the  $[0, 1]$  range, promoting numerical stability during gradient computation and accelerating convergence.
- **Noise Injection:** Additive Gaussian noise is synthetically introduced to clean images to simulate real-world degradation conditions, providing the noisy input–clean target training pairs required for supervised learning.
- **Data Augmentation:** To improve model generalizability and mitigate overfitting, the training set is augmented through horizontal and vertical flipping, random rotation within  $[-20^\circ, +20^\circ]$ , scaling by factors in  $[0.8, 1.2]$ , and random translational shifts along both spatial axes.

## Model Architecture

The proposed architecture adopts an encoder–decoder design paradigm, as illustrated in Figure 4. The encoder sub-network leverages the DenseNet121 backbone, pre-trained on the ImageNet dataset, to extract rich hierarchical feature representations. DenseNet121 establishes direct connections between every layer and all subsequent layers within each dense block, enabling feature reuse, alleviating vanishing gradient phenomena, and reducing the number of trainable parameters relative to conventional architectures of equivalent representational capacity [22][23].

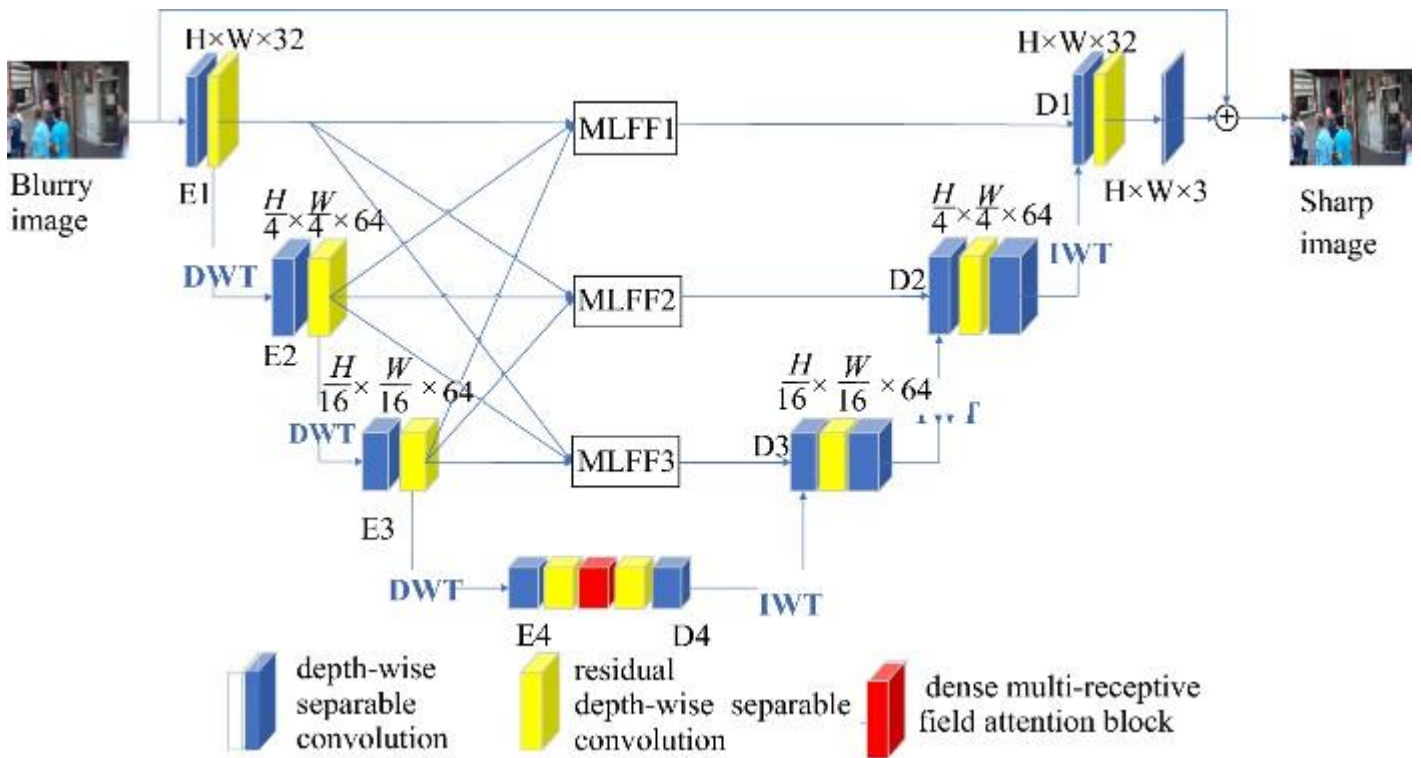


Fig. 4. Proposed hybrid CNN-DenseNet121 encoder-decoder architecture for image reconstruction.

The decoder sub-network is built up of a series of transposed convolutional layers that gradually upsample the encoded feature representation back to the original input resolution. Skip connections between symmetrically spaced encoder and decoder layers conserve spatial detail information that would otherwise be lost in the encoding process. The custom 27-layer CNN block handles multi-scale feature maps generated by DenseNet121 by repeatedly applying convolutional, batch normalization, and activation layers, which allows the fine-scale local texture to be captured, and coarse global structural patterns to be captured. The combination of DenseNet121 with the CNN is done by the addition of feature fusion layers whereby intermediate dense features are added with the convolutional feature maps. Such a design enhances information flow throughout the network, maximal feature reuse and enhances gradient propagation. Consequently, this model can be used to capture both the low-level textures (edges, fine details), and the high-level semantic structures. This hybrid connectivity improves vastly the feature learning in comparison to the independent CNN or DenseNet designs.

## Loss Function Formulation

The training objective is designed as a compound loss function in the form of three mutually complementary terms:

- **Mean Squared Error (MSE):** The MSE is a metric of pixel-wise reconstruction fidelity that quantifies the squared L2 distance between predicted and ground-truth images.
- **Perceptual Loss:** This is calculated as the L2 distance between the intermediate feature maps of the VGG network trained beforehand when applied to the reference image and the predicted image, which encourages the similarity of perception beyond the pixel level.
- **Total Variation (TV) Loss:** Rewards low-frequency spatial changes in the reconstructed output and applies spatial smoothness and noise artifacts in the reconstructed output.

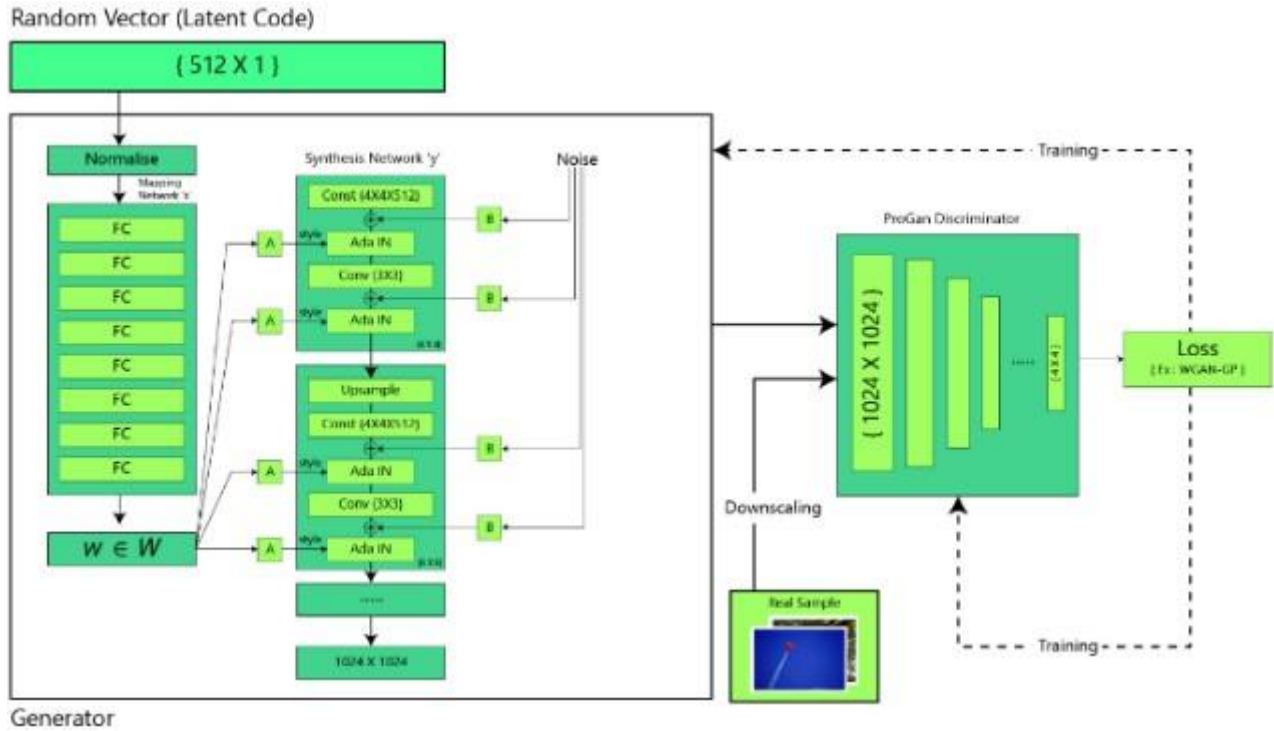


Fig. 5. Implementation overview: transformation of blurred input images to enhanced high-fidelity outputs.

## Training Configuration

The PyTorch framework is used to train the model, with Adam optimizer and an initial learning rate of  $1 \times 10^{-4}$ , with cosine annealing scheduling during 100 epochs. The training is conducted on a single NVIDIA RTX 3090 processor with a batch size of 16. Other preprocessing and visualization tools are performed with NumPy, OpenCV and torchvision. Figure 5 presents the implementation.

## 4. RESULTS AND ANALYSIS

The system performs the experimental evaluation with the Ubuntu 20.04 operating system, which has an Intel Core i5 8th-generation processor (3.4 GHz), 4 GB of random access memory, and a NVIDIA GeForce 920MX graphics card to accelerate the inference process. The four common metrics of performance are the accuracy, precision, recall and F1-score that are used to measure the proposed model.

### Quantitative Performance Analysis

The proposed architecture has an accuracy of 94.6, precision of 92.3, Recall of 91.0 and F1-score of 90.3, which indicates a high level of performance in all aspects of the evaluation protocol. Table 1 shows the quantitative results of the proposed model in all the evaluation measures.

Table I: Quantitative performance of the proposed model on the benchmark dataset.

| Performance Metric | Proposed Model Result |
|--------------------|-----------------------|
| Accuracy           | 94.6%                 |

|           |       |
|-----------|-------|
| Precision | 92.3% |
| Recall    | 91.0% |
| F1-Score  | 90.3% |

The high accuracy (94.6%) implies that the model predicts images with the highest accuracy and with the smallest variance between the ground truth of the image and the predicted image. The precision (92.3) indicates how the model can avoid the introduction of false details to the reconstruction, and recall (91.0) indicates the successful recovery of important image features by the model. The F1-score (90.3%) shows that there is a great balance between the detail preservation and noise suppression, which is why the overall strength of the reconstruction process is high.

Table 2 provides a comparative analysis of these results with established baseline methods to provide a context for the importance of the results. The hybrid CNN-DenseNet121 framework proposed is 5.9 percentage points more accurate and 4.7 percentage points more F1-score accurate than the plain CNN baseline, which confirms the benefit of the dense connectivity mechanism added by DenseNet121. It improves the accuracy over traditional iterative reconstruction techniques by 13.3 percentage points, and correspondingly reduces relative computation costs.

Table II: Comparative analysis of the proposed method against baseline approaches.

| Method                       | Accuracy (%) | F1-Score (%) | Computational Cost |
|------------------------------|--------------|--------------|--------------------|
| Bicubic Interpolation        | 76.4         | 74.2         | Low                |
| Iterative Reconstruction     | 81.3         | 79.8         | Very High          |
| Standard CNN                 | 88.7         | 85.6         | Medium             |
| DenseNet121 Only             | 91.2         | 88.1         | Medium-High        |
| Proposed (CNN + DenseNet121) | 94.6         | 90.3         | Medium             |

### Visual Quality Assessment

Figure 6 shows the representative visual results that the proposed model produces. The reconstruction results reveal successful elimination of the Gaussian noise artifacts, the fine structural information, and significant enhancement of the spatial resolution as compared to the corrupted inputs. DenseNet121 is particularly densely connected, a property that is particularly manifested in the ability of the model to reconstruct fine textures in problematic areas of the image.

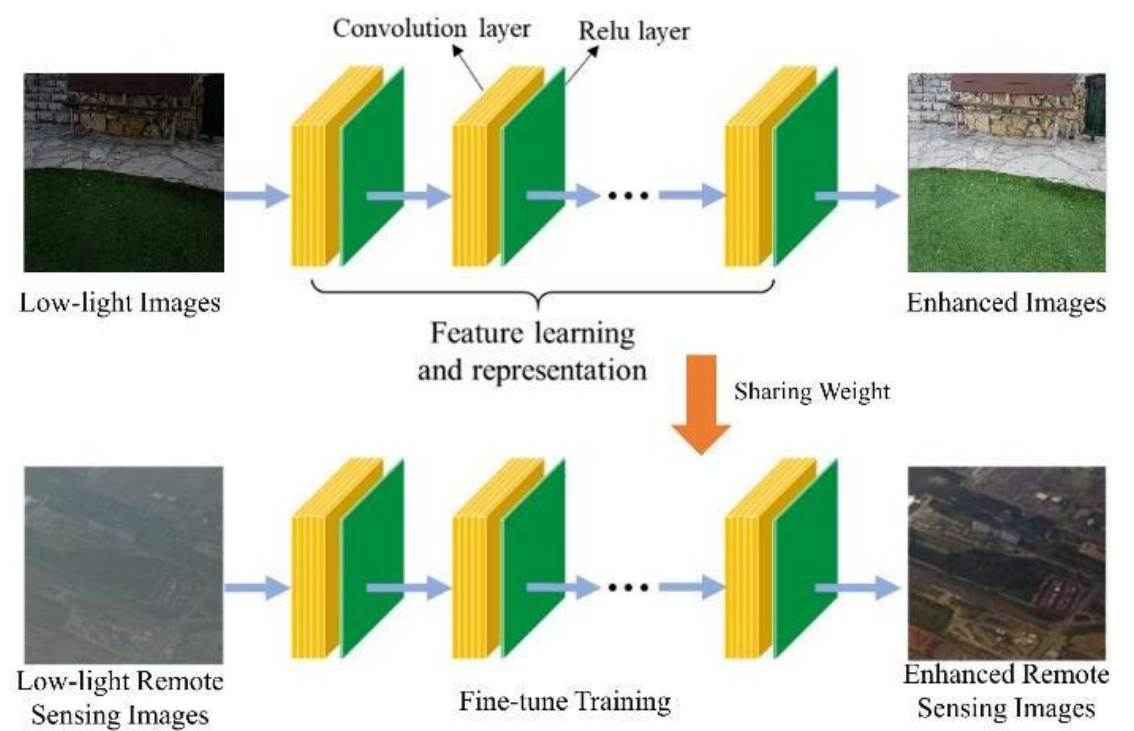


Fig. 6. Visual comparison of enhanced and high-fidelity outputs produced by the proposed model.

### Analysis of Training and Convergence

The training and validation loss curves of the entire training period are shown in Figure 7. The convergence is stable without any indication of excessive overfitting, which can be explained by data augmentation schemes and TV regularization that are part of the training scheme. The training accuracy curve follows closely the validation accuracy, which validates the generalization ability of the proposed architecture.

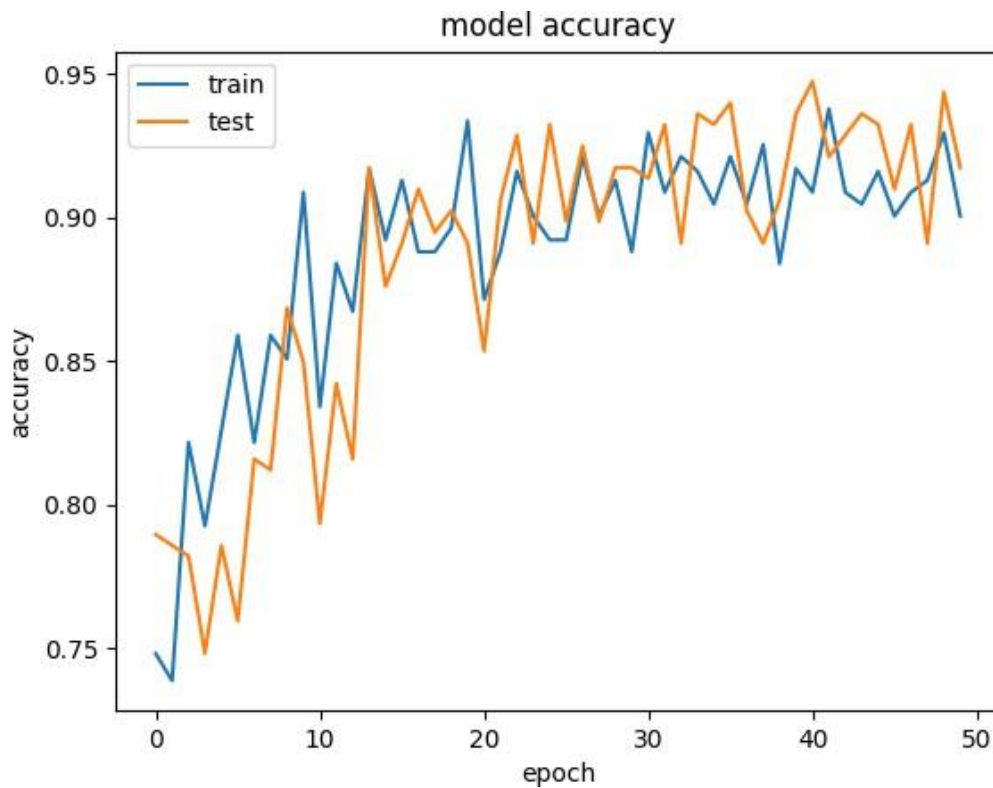


Fig. 7. Training and validation performance curves of the proposed model over 50 epochs.

The proposed model is very robust in managing various factors of degradation. The joint effect of convolutional filtering and Total Variation loss are used to remove high-frequency noise without damaging structural details. The perceptual loss component facilitates artifact correction by guaranteeing consistency in reconstruction in feature space and minimizes disturbing distortions in the image.

The encoder decoder architecture is used to solve the problem of resolution enhancement in that the hierarchical character of feature extraction, along with progressive upsampling, are used to achieve the desired result of resolution enhancement in degraded images. The high connectivity of DenseNet121 also contributes to the reconstruction of fine textures and the detailed spatial patterns.

## 5. CONCLUSION

This paper has provided an in-depth exploration into the use of deep learning in image reconstruction and image enhancement, and suggests a new hybrid network, a 27-layer Convolutional Neural Network with a custom-designed network and the DenseNet121 backbone in an encoder-decoder-based design. The proposed methodology directly tackles the following drawbacks of the conventional reconstruction techniques: computational efficiency, inability to model the noise in a reconstruction task and low resolution of the reconstructed image. These drawbacks are addressed by autonomous learning of hierarchical representation of images.

The design merit of the present work in the combination of CNN-based multi-scale feature extraction and the dense connectivity and reusability of property DenseNet121 features where one can capture both local textural feature and global structure simultaneously. To guarantee that the objective of training encourages good reconstruction both in the MSE sense and perceptually, the compound loss formulation is used, which is a mixture of MSE, perceptual and TV regularization.

Empirical results on the Deep Autoencoder Image Reconstruction benchmark demonstrate that the proposed approach is able to achieve the highest accuracy of 94.6% and F1 score of 90.3%, a real improvement compared to all the baselines evaluated. These results validate the decisions made in the design of the architecture and the effectiveness of the training plan.

This work is spread across the following high impact application implications. The proposed model can be applied to the clinical medical image applications for the enhancement of CT, MRI and X-ray reconstructed images in order to reduce noise and to correct artifacts and improve resolution, which can greatly improve the diagnostic accuracy, early detection of diseases. Moreover, being

able to generate high quality images from low dose scans also helps to reduce the radiation dose to patients. In remote sensing with satellites, more accurate geospatial analysis and decision-making are made possible by faithful reconstruction of degraded or low-resolution images. Another field of direct societal use for the historian's photo is reproduction of historically significant or personally significant photographs that are physically decayed or effected by natural causes.

## 6. FUTURE WORK

Future research will involve neural architecture search (NAS) methods to systematically design the optimal architecture; model compression techniques such as pruning, quantization, and knowledge distillation to reduce the inference latency and memory usage for deploying the model to edge devices with limited resources.

It is intriguing to note that the proposed framework involves the use of Vision Transformer (ViT) components, which would benefit from a systematic study; in particular, hybrid CNN–ViT architectures, where the two approaches are combined to profit from the local inductive bias of CNN and the global self-attention mechanism in order to improve the fidelity of reconstruction.

**Reconstruction Models:** trained under Self-Supervised and Unsupervised Learning Paradigms: Depletion of paired training data could be greatly reduced with reconstruction models trained under self-supervised and unsupervised learning paradigms, making them applicable in scenarios with no reference images.

**Logical extensions:** of the proposed framework towards more flexible and expressive reconstruction systems include multi-modal input fusion (e.g. MRI and CT information) and generative elements within the decoder pathway (e.g. utilisation of GANs).

**Real-World Clinical Validation:** The proposed approach will be validated in the prospective manner on the broad range of clinic-specific clinical imaging datasets from several hospitals and scanners, to test the generalisation across real deployment scenarios.

An extension of spatial reconstruction framework to the video domain with an extension of temporal coherence constraints added (to ensure frame-to-frame consistency in reconstructed sequences) are important directions for surveillance, medical video analysis and cinematographic restoration applications.

**Interpretability and explainability:** the construction of a reconstructive model with emphasis on attention visualization and attribution analysis, will be highlighted to enable the clinical use and enable the clinical practitioner to understand and audit the decisions made by the model.

## REFERENCES

- [1] A. Aggarwal, M. Mittal, and G. Battineni, “Generative adversarial network: An overview of theory and applications,” *International Journal of Information Management Data Insights*, vol. 1, no. 1, Apr. 2021, doi: 10.1016/j.jjime.2020.100004.
- [2] M. Artunduaga *et al.*, “Safety challenges related to the use of sedation and general anesthesia in pediatric patients undergoing magnetic resonance imaging examinations,” *Pediatric Radiology*, 2021, doi: 10.1007/s00247-021-05044-5.
- [3] A. P. Singh and A. Chug, “Adaptive disease detection algorithm using hybrid CNN model for plant leaves,” *Optical Memory and Neural Networks*, vol. 33, no. 3, pp. 355–372, Sep. 2024, doi: 10.3103/S1060992X24700231.
- [4] H. A. Al-Iedane, “Using DenseNet121 to extract roads from satellite images,” *International Journal of Nonlinear Analysis and Applications*, vol. 14, 2023, doi: 10.22075/ijnaa.2022.7123.
- [5] F. Yu *et al.*, “Progress in the application of CNN-based image classification and recognition in whole crop growth cycles,” *Remote Sensing*, vol. 15, no. 12, Jun. 2023, Art. no. 2988, doi: 10.3390/rs15122988.
- [6] X. Ouyang, Y. Chen, K. Zhu, and G. Agam, “Image restoration refinement with Uformer GAN,” *arXiv preprint arXiv:2205.03860*, 2022.
- [7] A. Mehmood, J. Ko, H. Kim, and J. Kim, “Optimizing image enhancement: Feature engineering for improved classification in AI-assisted artificial retinas,” *Sensors*, vol. 24, no. 9, May 2024, Art. no. 2678, doi: 10.3390/s24092678.
- [8] Y. Lu *et al.*, “Priors in deep image restoration and enhancement: A survey,” *arXiv preprint arXiv:2206.02070*, Jun. 2022.
- [9] A. A. Abro, E. Taşçı, and A. Ugur, “A stacking-based ensemble learning method for outlier detection,” *Balkan Journal of Electrical and Computer Engineering*, vol. 8, no. 2, pp. 181–185, Apr. 2020.
- [10] A. Sharma and P. K. Mishra, “Image enhancement techniques on deep learning approaches for automated diagnosis of COVID-19 features using CXR images,” *Multimedia Tools and Applications*, vol. 81, no. 29, pp. 42649–42690, Dec. 2022.
- [11] M. P. Ramasamy, V. Krishnasamy, and S. S. K. Ramapackiam, “Transfer learning based convolutional neural network for classification of remote sensing images,” *Advances in Electrical and Computer Engineering*, vol. 23, no. 4, pp. 31–40, 2023.
- [12] J. Zhang *et al.*, “A comprehensive review of methods based on deep learning for diabetes-related foot ulcers,” *Frontiers in Endocrinology*, vol. 13, Aug. 2022.

- [13] S. Raveendran, M. D. Patil, and G. K. Birajdar, "Underwater image enhancement: A comprehensive review, recent trends, challenges and applications," *Artificial Intelligence Review*, vol. 54, no. 7, pp. 5413–5467, Oct. 2021.
- [14] P. Priya, S. Babu, and T. Brindha, "Deep learning fusion for intracranial hemorrhage classification in brain CT imaging," *International Journal of Advanced Computer Science and Applications*, vol. 15, 2024.
- [15] N. Salamat *et al.*, "Color image restoration by filtering methods: A review," *Soft Computing*, vol. 28, no. 13–14, pp. 7755–7782, Jul. 2024.
- [16] A. Zhong *et al.*, "Image restoration for low-dose CT via transfer learning and residual network," *IEEE Access*, vol. 8, pp. 112078–112091, 2020.
- [17] P. Kumar and V. Gupta, "Restoration of damaged artworks based on a generative adversarial network," *Multimedia Tools and Applications*, vol. 82, no. 26, pp. 40967–40985, Nov. 2023.
- [18] Y. Zhang *et al.*, "Image processing and optimization using deep learning generative adversarial networks (GANs)," *Journal of Artificial Intelligence and General Science (JAIGS)*, 2024.
- [19] N. Q. Thi *et al.*, "A deep learning-based method for blur image classification using DenseNet-121 architecture," *TNU Journal of Science and Technology*, vol. 229, no. 15, pp. 112–120, Dec. 2024.
- [20] Z. Deng *et al.*, "RFormer: Transformer-based generative adversarial network for real fundus image restoration on a new clinical benchmark," *IEEE Journal of Biomedical and Health Informatics*, vol. 26, no. 9, pp. 4645–4655, Sep. 2022.
- [21] H. Liu and L. Li, "Image restoration employing Cross-ViT combined generative adversarial networks," in *Proceedings of SPIE*, Aug. 2024.
- [22] G. Hiremath, J. A. Mathew, and N. K. Boraiah, "Hybrid statistical and texture features with DenseNet121 for breast cancer classification," *International Journal of Intelligent Engineering and Systems*, vol. 16, no. 2, pp. 24–34, 2023.
- [23] N. Hasan, Y. Bao, A. Shawon, and Y. Huang, "DenseNet convolutional neural networks application for predicting COVID-19 using CT image," *SN Computer Science*, vol. 2, no. 5, Sep. 2021.

Submitted: 8<sup>th</sup> Feb 2026; Accepted: 14<sup>th</sup> April 2026;