

AI-Driven Error Correction for Real-Time Wireless Image Transmission: Integrating Banach Space Principles into Deep Learning

Charles Kinyua Gitonga^{1,*}, Ismael Munene Kwenga², and Sammy Musundi³

ABSTRACT

Wireless communication has been widely adopted owing to its flexibility. However, it introduces many limitations, such as noise interference for real-time image transmission, which is key to modern applications, such as telemedicine. This study seeks to integrate Banach space principles with deep learning to improve the error correction in wireless image transmission. This research is justified by the need to maintain high-quality images in real time, despite unpredictable wireless channel errors. They proposed a convolutional autoencoder enhanced with an iterative refinement module that enforces contraction mappings via spectral normalization and contraction regularization. The model was trained on the CIFAR-10 dataset using noise simulation and advanced data augmentation, and evaluated using Peak Signal to Noise Ratio (PSNR), Structural Similarity Index (SSIM), and inference time metrics. The experiments including dynamic lambda scheduling, demonstrated that under moderate Gaussian noise the model achieves up to 23.66 dB PSNR and 0.8489 SSIM while processing over 600 frames per second. Ethical considerations were addressed using publicly available data, and the code and methodology are well documented for reproducibility. Study limitations included the use of low-resolution images and simulated noise, which may not fully capture real-world challenging conditions. Future work will extend these results to larger and more complex datasets and real transmission scenarios.

Submitted: May 31, 2025

Published: October 10, 2025

 10.24018/ejai.2025.4.5.71

¹ Computer Science, Chuka University, Kenya and Tharaka University, Kenya.

² Computer Science, Tharaka University, Kenya.

³ Pyscial Science, Chuka University, Kenya.

* Corresponding Author:

e-mail: cgkinyua@chuka.ac.ke

Keywords: Banach space, contraction mapping, iterative refinement, Wireless image transmission.

1. INTRODUCTION

In recent years, wireless communication has been widely deployed across various sectors. Real-time wireless image transmission is a critical component in modern applications such as telemedicine, surveillance, and interactive streaming. The over-the-air channels used in wireless communication inherently introduce challenges, such as interference, fading, and packet loss, which compromise the integrity of real-time image transmission [1].

Wireless channels suffer from errors and packet losses due to interference, fading, and dynamic channel conditions. These factors significantly degrade the image quality and user experience [2]. Traditional approaches in wireless communication are primarily based on separate source and channel coding techniques, such as Low-Density Parity-Check (LDPC) and Reed Solomon codes. These approaches perform well in stable environments, but fall

short when exposed to dynamic and harsh transmission conditions [3].

Traditional error-correcting codes, including Reed–Solomon and low-density parity-check (LDPC) codes, often introduce latency and may fail under unpredictable conditions [3]. Consequently, there is an increasing demand for adaptive and robust methods that can maintain real-time image quality. Deep learning has emerged as a promising approach to tackle these challenges by offering the ability to learn complex noise patterns and efficiently restore images [4], [5]. Joint source-channel coding (JSCC), enhanced by deep learning, enables the simultaneous optimization of compression and error resilience, significantly improving performance in noisy environments [5].

Integrating the Banach space principles into deep learning provides a theoretical foundation for stability and convergence during iterative error correction [6]. The



application of contraction mapping through spectral normalization and contraction regularization ensures that the error-correction outputs are reliable and consistent. This study bridges mathematical theorems with advanced neural architectures to address this critical gap in achieving reliable real-time image transmission.

The Banach space theory further strengthens the proposed approach by ensuring convergence in iterative error-correction processes. In particular, contraction mapping guarantees that the outputs remain bounded and stable, even when subjected to substantial noise perturbations. This mathematical foundation addresses a critical limitation of purely data-driven methods, which may diverge under extreme conditions [6]. With advancements in GPU hardware, such as Nvidia A30, real-time processing of complex models has become feasible without significant latency [7].

Despite these advancements, several challenges remain. The following questions form the core motivation of this study. They guided both the theoretical analysis and the experimental validation of this research. 1) How can robust training be ensured across diverse wireless conditions? 2) Can Banach space properties universally guarantee convergence beyond the controlled simulations? 3) How does the proposed framework compare to state-of-the-art error-correction techniques?

1.1. Problem Statement

Real-time wireless image transmission often encounters unpredictable errors due to fading, interference, and limited bandwidth. Conventional error-correcting codes, while effective under stable conditions, struggle to adapt to rapid channel fluctuations and impose computational overhead that may lead to latency. For time-critical applications, such as telemedicine and interactive streaming, these limitations result in degraded image quality or delays in data transmission. Therefore, there is an urgent need for a robust, low-latency error-correction technique that adapts to changing channel states while ensuring stable image recovery.

This study proposes the integration of deep learning with Banach space principles to address these challenges. The goal was to develop an end-to-end architecture capable of minimizing error artifacts, maintaining real-time performance on modern GPU hardware, and verifying convergence and reliability across various transmission scenarios. In addition, evaluating the computational efficiency of different hardware configurations will ensure the feasibility of practical deployment.

1.2. Research Objectives

The objectives that this research aims to achieve are to;

1. Develop a deep learning framework for real-time wireless image transmission that integrates joint source-channel coding with an AI-driven decoder.
2. Incorporate Banach space properties to ensure iterative stability and convergence during error correction.
3. Evaluate the proposed method on a GPU-based setup (e.g., Nvidia A30) to confirm its low latency and high fidelity under simulated noisy conditions.

4. Benchmarking the proposed framework against classical error-correction strategies and demonstrating improvements in image quality and transmission reliability.
5. The computational efficiency of the model was assessed using various hardware configurations to understand its scalability and practical application.

1.3. Significance of the Study

This paper presents a novel deep-learning approach for real-time wireless image transmission. This approach integrates Banach space contraction principles with iterative refinement combined with dynamic regularization. The proposed autoencoder-based model demonstrated significant improvements in reconstruction accuracy and perceptual quality under challenging noise conditions. The study simulated real-world noise through advanced data augmentation and stabilizing training using contraction-based regularization.

The model achieved higher PSNR and SSIM scores while maintaining ultrafast inference times suitable for real-time applications. The integration of mathematical rigor with a practical neural network design offers a scalable, efficient, and robust framework for noise-resilient image transmission in modern communication environments.

The rest of this paper is organized as follows. Section 2 presents related works in the literature. Section 3 discusses the methodology of this study. Section 4 presents and discusses the results. Finally, Section 5 concludes the study with recommendations in Section 6.

2. LITERATURE REVIEW

In this section, we present a review of literature related to this research.

2.1. Historical Context of Wireless Image Transmission

Wireless image transmission has evolved significantly from the early days of mobile communications, when bandwidth limitations and rudimentary transmission techniques restricted multimedia content to low-resolution and text-based applications [1]. Early cellular networks, such as 2G and 3G, primarily facilitated text messaging and minimal multimedia support through General Packet Radio Service (GPRS) and Enhanced Data Rates for GSM Evolution (EDGE) [8]. As mobile technology progressed, fourth-generation (4G) networks introduced orthogonal frequency-division multiple access (OFDMA).

OFDMA significantly boosts the data throughput for real-time video and image transmission [9]. The advent of 5G has further pushed boundaries. It offers ultra-low latency and high bandwidth, which makes high-fidelity real-time image transmission a practical reality [10].

Despite these advancements, wireless image transmission continues to face challenges related to dynamic channel conditions, fading, and packet loss, which severely affect image quality [11]. Traditional error correction methods such as Reed Solomon and Low-Density Parity-Check (LDPC) codes offer protection against errors, but

often fail under rapid channel variations or high-resolution data [12]. Emerging AI-driven methods are increasingly being considered to address these challenges by leveraging deep learning to adaptively correct errors and maintain transmission quality [13].

2.2. Traditional Error Correction Techniques

Traditional error-correction methods play a fundamental role in ensuring data reliability in wireless communication. Forward Error Correction (FEC) techniques, such as Reed–Solomon and Low-Density Parity-Check (LDPC) codes, are widely used to correct data blocks without the need for retransmissions, employing redundancy and iterative decoding to recover lost information [12]. Automatic Repeat Request (ARQ) protocols and Hybrid ARQ schemes are built on this using feedback mechanisms to selectively retransmit erroneous packets. In particular, hybrid ARQ attempts to balance robustness and latency by combining FEC with selective retransmission [14]. In addition to these communication oriented approaches, mathematical frameworks also provide essential foundations for stability and convergence in error-correction processes. Functional analysis, particularly Banach and Sobolev space theory, offers rigorous tools for understanding iterative stability in complex systems [15]. At the same time, objective image quality assessment has advanced through the development of the Structural Similarity Index (SSIM), which better aligns with human perception compared to earlier pixel-based metrics [16]. These theoretical and perceptual insights complement communication strategies, creating a broader foundation for reliable real-time image transmission.

Despite the success of traditional error-correction mechanisms in conventional wireless systems, these classical approaches exhibit limitations when applied to high-resolution or real-time image transmissions. FEC and ARQ systems generally assume static or predictable error environments and lack the adaptability to cope with rapidly changing channel conditions or burst errors. Furthermore, they treat all parts of an image equally, failing to prioritize the visually critical regions. This results in increased latency, degraded image quality, and inconsistent performance, particularly in bandwidth.

2.3. Deep Learning Approaches for Wireless Image Transmission

Deep learning has introduced flexible and adaptive methods for wireless image transmission with notable improvements over traditional coding schemes. A prominent approach is joint source–channel coding (JSCC), which integrates compression and error correction into a unified neural framework. By encoding images into robust latent representations, JSCC models maintain image fidelity even under adverse transmission conditions, often outperforming classical separated coding strategies [17], [18].

Convolutional Autoencoders (CAEs) further enhance resilience by learning to reconstruct distorted images through the supervised training of noisy inputs. Their architecture efficiently preserves spatial structures, making

them suitable for real-time scenarios in which low latency and robustness are essential [18].

Generative Adversarial Networks (GANs) offer another layer of refinement, particularly to enhance perceptual quality. Through adversarial training, GANs produce more realistic reconstructions with sharper details, particularly in texture-rich regions. These capabilities make GAN-based methods ideal for applications requiring high-quality visual outputs, despite significant transmission errors [19].

2.4. Banach Space Concepts in Convergence and Stability

The Banach space theory offers a rigorous mathematical foundation for ensuring convergence and stability in iterative learning and error-correction processes. The core of this theory lies in the Banach Fixed-Point Theorem, which asserts that contraction mappings in a complete metric space will converge to a unique fixed point. When applied to deep learning, these contraction principles guarantee that iterative outputs remain stable, even under significant noise [20].

Incorporating Banach space principles into neural network architectures enhances the reliability of real-time error corrections. This ensures that the reconstructions do not diverge during training or inference. By embedding contraction constraints into layers or refinement modules, networks can maintain predictable and robust behavior, which is an essential characteristic for deployment in volatile wireless environments [20].

2.5. Benchmarking and Performance Metrics

Robust benchmarking and performance evaluation are crucial for assessing the effectiveness of wireless image-transmission systems, particularly those enhanced by deep learning. Commonly used quantitative metrics include Peak Signal-to-Noise Ratio (PSNR) and Structural Similarity Index Measure (SSIM) [4], [21]. PSNR assesses pixel-level fidelity by comparing the reconstructed image to the original in terms of logarithmic decibel values, whereas SSIM evaluates perceptual quality by measuring structural and luminance similarities between the two images [18], [19]. These metrics are widely accepted owing to their ease of computation and correlation with visual distortions.

However, PSNR and SSIM alone may not fully capture the visual realism of the reconstructed images, particularly in human-centric applications. To address this limitation, recent studies have proposed additional perceptual quality metrics, such as the Learned Perceptual Image Patch Similarity (LPIPS) and the Mean Opinion Score (MOS), which better align with subjective visual assessment [19], [20]. These emerging metrics help bridge the gap between objective evaluation and user experience.

In addition to image quality, real-time performance metrics play a critical role in determining the feasibility of deployment of transmission systems. Key indicators include inference latency, measured in milliseconds per frame, and throughput, typically expressed as frames per second (FPS). High FPS and low latency are essential for applications such as live video streaming, telemedicine, and

interactive surveillance, where delays can compromise the user experience or system reliability [21].

Energy efficiency and memory usage are also increasingly important benchmarks, especially in edge computing contexts, where devices operate under strict resource constraints. Evaluating power consumption, model size, and inference efficiency allows researchers to tailor models for real-world deployment on embedded or mobile hardware [21]. However, this study was limited by the PSNR and SSIM metrics. Future research would therefore build on the findings of this research to extend it to additional perceptual quality metrics.

2.6. Gaps in the Literature and Contribution of this Study

Although deep learning has advanced wireless image transmission, significant gaps remain, including limited scalability, lack of real-world generalizability, static regularization, weak perceptual quality assessment, rigid architectures, and insufficient focus on energy efficiency. This study has several limitations. We introduced a scalable and modular deep learning framework designed for real-time performance, achieving a 614.47 FPS. Although tested on CIFAR-10, the architecture supports future adaptation to higher-resolution data. To enhance realism, the model was evaluated under varied noise conditions beyond typical static simulations using Gaussian and burst noise to reflect diverse wireless distortions.

A key strategy for the proposed method is the adoption of dynamic contraction regularization. This improved stability and adaptability by varying the regularization strength (λ) during training. This contributed to more robust reconstructions compared with the fixed regularization strategies. The study also considers computational efficiency by optimizing the inference time and model simplicity. This is recommended in constrained environments. These contributions collectively bridge theoretical and applied research by offering a foundation for scalable, stable, and practical wireless image-restoration systems.

3. METHODOLOGY

In this section, we present the research methodology adopted in this study.

3.1. Research Design and Experimental Setup

We adopted an experimental research design to evaluate the effectiveness of integrating the Banach space principles into deep learning models for error correction in real-time wireless image transmission. This study aims to formulate the problem of image denoising, focusing on reconstructing clean images from artificially corrupted versions that simulate real-world wireless transmission errors. The experiments involved artificially corrupting CIFAR-10 dataset images to mimic real-world wireless transmission errors. We trained the models to restore the original clean images by minimizing the reconstruction errors. This experimental setup allowed us to evaluate the model performance under varying noise conditions, such as Gaussian and salt-and-pepper noise.

TABLE I: PYTHON SCRIPTS

Script name	Purpose
Autoencoder simulation.py	Initial training of the convolutional autoencoder model, establishing baseline performance.
Autoencoder refinement.py	Enhanced autoencoder with iterative refinement to improve PSNR and SSIM metrics.
Autoencoder refinement v2.py	Further refined model with explicit validation splits to prevent overfitting.
Hyperparameter tuning.py	Automated hyperparameter testing, saving results to CSV for analysis.
Evaluate test set.py	Evaluation on CIFAR-10 test set, generating PSNR, SSIM, and inference time metrics.
Ablation and data augmentation.py	Ablation study comparing baseline versus refined models and augmentation effects.
Banach integration experiments.py	Tests Banach space principles with contraction regularization and stability analysis.
Banach comparison analysis.py	Comparative analysis of models trained with and without Banach integration.
Extended evaluation.py	Detailed testing under varied noise conditions, generating comprehensive robustness metrics.
Research extended pipeline.py	Complete integrated research pipeline for dynamic contraction regularization and visualization.

3.2. Python Code Structure and Organization

To ensure reproducibility, modularity, and maintainability, the research implementation was organized through several Python scripts, each serving a specific role within the pipeline, as shown in Table I.

3.3. Dataset and Preprocessing

The CIFAR-10 dataset was used in this study. It consists of 60,000 labeled color images, each with dimensions of 32×32 pixels. The dataset was systematically divided into three subsets: the Training Set (45,000 images), Validation Set (5000 images), and Test Set (10,000 images). This division allows for robust model training and evaluation.

To enhance the ability of the model to generalize to diverse conditions, various data augmentation techniques have been employed. Basic augmentations include converting images into tensor representations and normalizing their values. Extended augmentation techniques include horizontal flipping, random cropping, rotation, and color jittering. These augmentations increase the diversity of the training data and improve the robustness of the model against unseen noise patterns.

To simulate real-world wireless transmission errors, we introduced Gaussian noise and salt-and-pepper noise of varying intensities during training. Gaussian noise was simulated with random signal-to-noise ratios, whereas salt-and-pepper noise mimicked burst errors by randomly flipping pixel values, as shown in Fig. 1. This was performed to ensure that the model could adapt to different noise scenarios and remain robust across varying conditions.

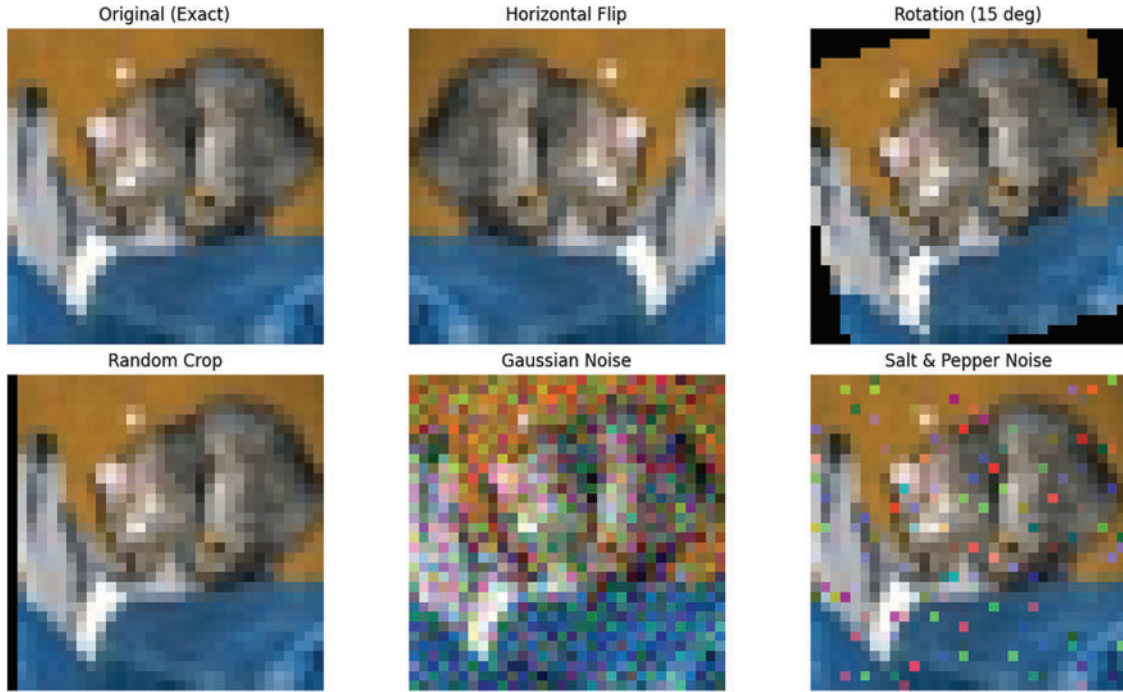


Fig. 1. Gaussian noise simulated.

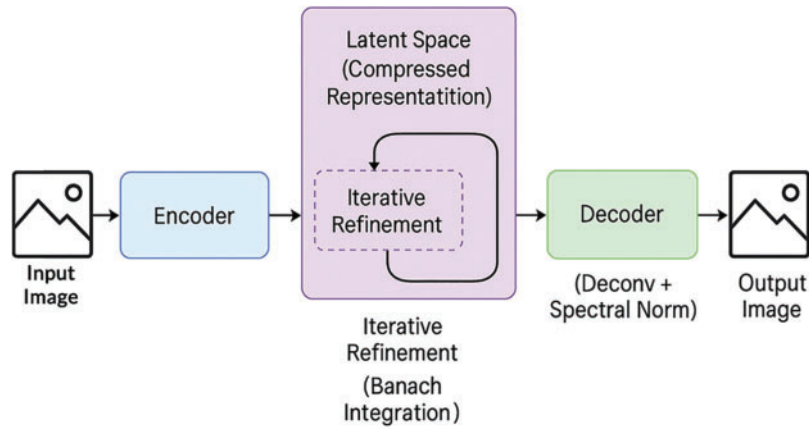


Fig. 2. Contraction mappings, source (author).

3.4. Model Architecture and Banach Space Integration

The core of the methodology involved designing a deep convolutional autoencoder (CAE). This model was chosen for its ability to encode noisy images into compressed latent representations while preserving essential features. The encoder network performs dimensionality reduction, whereas the decoder network reconstructs the original image from the latent representations.

To enhance stability, we incorporated Banach space principles by enforcing contraction mappings using the Banach Fixed-Point Theorem. This was achieved by embedding iterative refinement modules in the model architecture. Each iteration in the iterative refinement aims to progressively improve the image reconstruction quality, as shown in Fig. 2.

The iterative modules helped stabilize the output of the model by reducing the risk of divergent reconstructions. Spectral normalization was applied to maintain a constrained Lipschitz constant in the convolutional layers, ensuring that minor variations in the input would not

disproportionately affect the output. This regularization technique is critical for maintaining stable and predictable reconstruction performance.

3.5. Model Training and Hyperparameters

Training the model involved the use of the Adam optimizer, which was chosen for its adaptability and efficiency in handling large datasets. The model was trained with the following hyperparameters to leverage GPU capabilities: Learning Rate: 1×10^{-3} , Batch Size: 64, epoch: 100, and Precision; Mixed Precision (FP16).

The dynamic contraction regularization strategy involved progressively increasing the regularization parameter λ from 0.01 to 0.1 over the training epochs. This gradual increase allowed the model to strike a balance between the reconstruction quality and contraction stability.

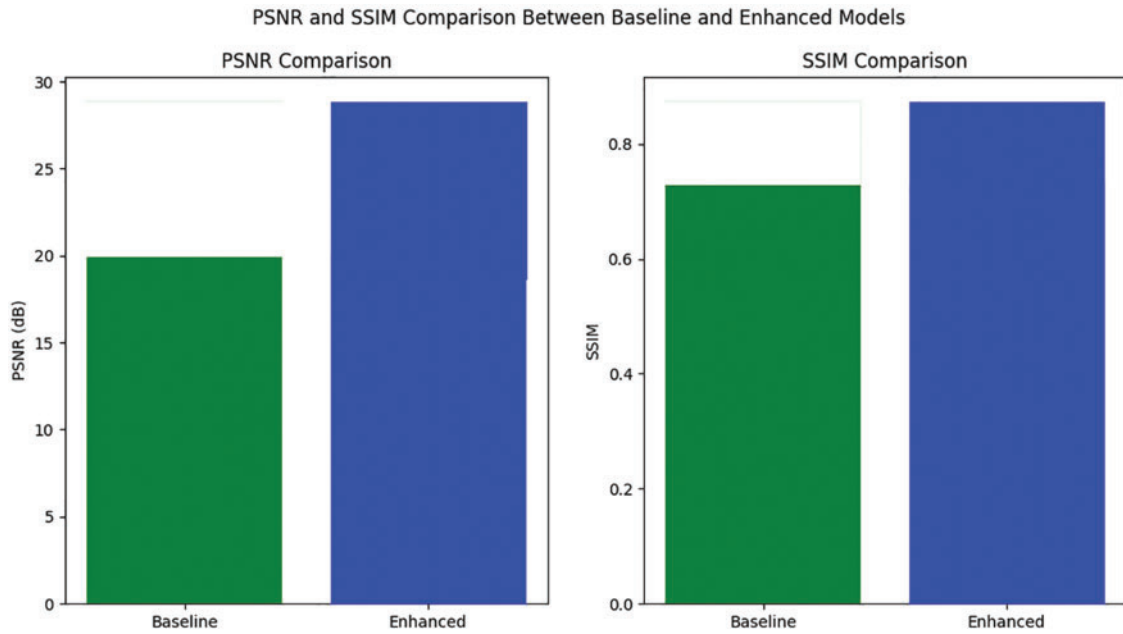


Fig. 3. Models comparison on PSNR and SSIM.

4. RESULTS AND DISCUSSION

In this section, we present experimental results and provide a detailed analysis. The methodology compares baseline models with refined models incorporating contraction mappings, dynamic contraction regularization, iterative refinement, and advanced data augmentation.

Fig. 3 shows a comparison of the Peak Signal-to-Noise Ratio (PSNR) and Structural Similarity Index Measure (SSIM) between the baseline and enhanced (proposed) models. The left subplot illustrates the PSNR values, showing that the enhanced model (blue bar) achieves a significantly higher PSNR (29 dB) than the baseline model (green bar, 20 dB), indicating better pixel-level fidelity. The right subplot displays the SSIM values, where the enhanced model again outperforms the baseline, achieving a value above 0.85, reflecting improved structural and perceptual similarity. These metrics collectively demonstrate the superiority of the enhanced model in preserving the image quality during wireless transmission.

To statistically validate the model's improvements, paired t-tests were conducted between the baseline and the PSNR and SSIM scores of the proposed model. The results indicated statistically significant improvements ($p < 0.05$), confirming the effectiveness of the model in enhancing the image quality under noisy conditions. The finalized autoencoder architecture featured symmetric convolutional encoding-decoding layers with batch normalization. Contraction regularization was applied in both the latent-space and output layers to promote robust feature representation and maintain consistent reconstruction accuracy.

4.1. Baseline Autoencoder Performance

Initially, we trained a baseline convolutional autoencoder without Banach integration or iterative refinement. The baseline model achieved a Peak Signal-to-Noise Ratio (PSNR) of 20.85 dB and a Structural Similarity Index Measure (SSIM) of 0.7133 on the CIFAR-10 validation

TABLE II: BASELINE MODEL PERFORMANCE METRICS

Metric	Value
PSNR (dB)	20.85
SSIM	0.7133
Inference time (ms)	0.63

TABLE III: IMPROVED PERFORMANCE THROUGH ITERATIVE REFINEMENT

Metric	Value
PSNR (dB)	28.83
SSIM	0.873
Inference time (ms)	0.02

set. The average inference time per image was 0.63 ms, confirming the feasibility of real-time processing. These baseline results provide a reference point for measuring the improvements offered by the proposed methods, as shown in Table II.

4.2. Evaluation and Testing Framework

The model evaluation focused on the CIFAR-10 test set. We measured the Peak Signal-to-Noise Ratio (PSNR) to assess pixel-level reconstruction fidelity, Structural Similarity Index Measure (SSIM) to quantify perceptual similarity relative to human vision, and Inference Time per image to verify real-time performance requirements. The model was subjected to diverse noise conditions, including Gaussian and burst noises, to assess its robustness. The performance metrics were calculated and visualized to effectively compare the baseline and enhanced (proposed) models, as shown in Fig. 3.

4.3. Improved Performance through Iterative Refinement

The integration of iterative refinement modules into the autoencoder architecture led to measurable improvements in the reconstruction quality. The enhanced model was trained for 20 epochs with early stopping based on the

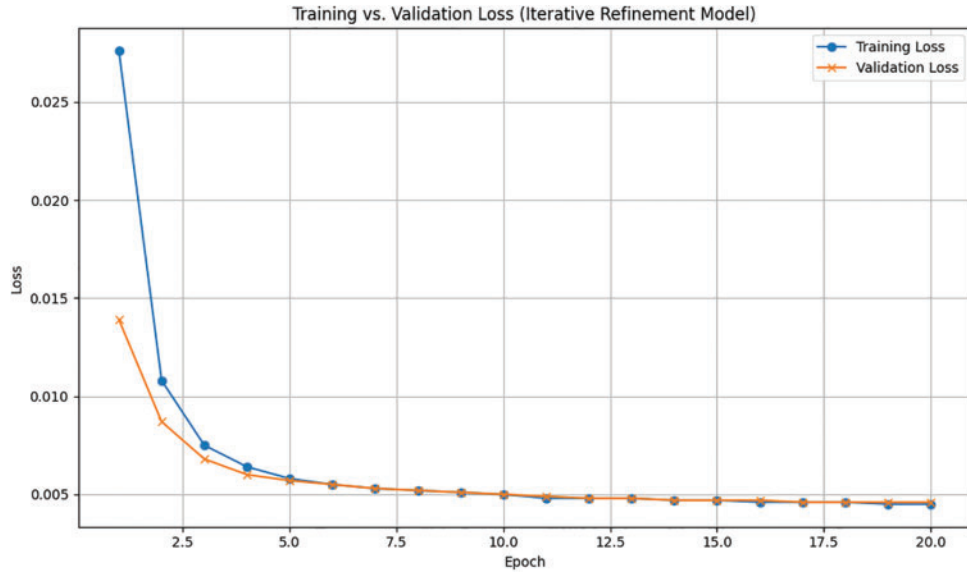


Fig. 4. Training vs. validation loss curves for iterative refinement model.

validation loss. The training loss consistently decreased, reaching 0.0045, whereas the validation loss stabilized at 0.0046, indicating convergence without overfitting. Performance on the CIFAR-10 test set under standard Gaussian noise conditions showed clear gains of PSNR of 28.83 db, SSIM of 0.873 and Inference time of 0.02 ms as shown on Table III.

These improvements demonstrate the clear advantage of iterative refinement of the proposed model in reducing reconstruction errors and enhancing visual quality.

Fig. 4 depicts the training and validation loss curves over 20 epochs for the iterative refinement model. The training loss initially decreased sharply, indicating rapid learning during the early stages. After approximately five epochs, the reduction in the training loss stabilized and converged steadily. This signaled an effective training without overfitting. The validation loss also followed the trend closely, quickly decreasing and leveling off alongside the training loss. This confirms the generalization of the model. The small gap between the training and validation curves throughout the training period suggests that the iterative refinement model effectively balances the learning accuracy and generalization performance. This demonstrates the stability and robustness of the learning image reconstruction tasks.

4.4. Ablation and Data Augmentation Studies

Ablation experiments were conducted to evaluate the individual and combined effects of iterative refinement and data augmentation strategies on the model performance. The baseline model, which lacked refinement and augmentation, achieved a PSNR of 20.24 dB and SSIM of 0.7016. Introducing iterative refinements alone resulted in improved fidelity, increasing the PSNR to 24.01 dB and the SSIM to 0.8112, confirming the efficacy of the iterative correction mechanism in reducing noise artifacts and enhancing structural integrity.

Interestingly, when basic data augmentation techniques were applied, such as horizontal flips and random crops,

the results showed no significant gain, with a slight reduction in SSIM to 0.7989 and PSNR dropping to 23.23 dB. This indicates that basic augmentation may not sufficiently enhance the robustness of the model against wireless transmission distortions. However, when extended augmentation strategies were introduced, including color jittering, rotation, and noise injection, perceptual quality improved significantly. While the PSNR slightly dropped to 23.02 dB, the SSIM increased to 0.8231, the highest among all configurations. This trade-off illustrates that extended augmentation enhances perceptual consistency and visual quality, even if the pixel-level similarity (as measured by PSNR) does not improve proportionally. Fig. 5 shows a performance comparison.

The refined model with iterative refinements achieves the highest PSNR, indicating improved numerical fidelity. In contrast, the refined model with extended augmentation yields the best SSIM, reflecting higher perceptual similarity. This shows that refinements and augmentations impact fidelity and perceptual quality differently.

4.5. Banach Space Integration Contraction Regularization

To evaluate the impact of Banach space principles on the robustness and convergence of the model, we implemented contraction regularization using different values of the regularization strength (λ). The performance results revealed a strong dependence on λ , highlighting the importance of balancing stability and expressiveness. When a relatively high contraction strength was applied ($\lambda = 0.10$), the model achieved the best performance with a PSNR 23.30 dB and SSIM of 0.8038. This result demonstrates that enforcing a meaningful contraction level encourages convergence to more stable and accurate reconstructions.

On the other hand, reducing λ to 0.05 led to a significant performance drop (PSNR = 19.07 dB, SSIM = 0.6692), indicating that excessive regularization constrained the model too tightly, impairing its learning capacity. Conversely, with a lower λ of 0.01, the model gained some

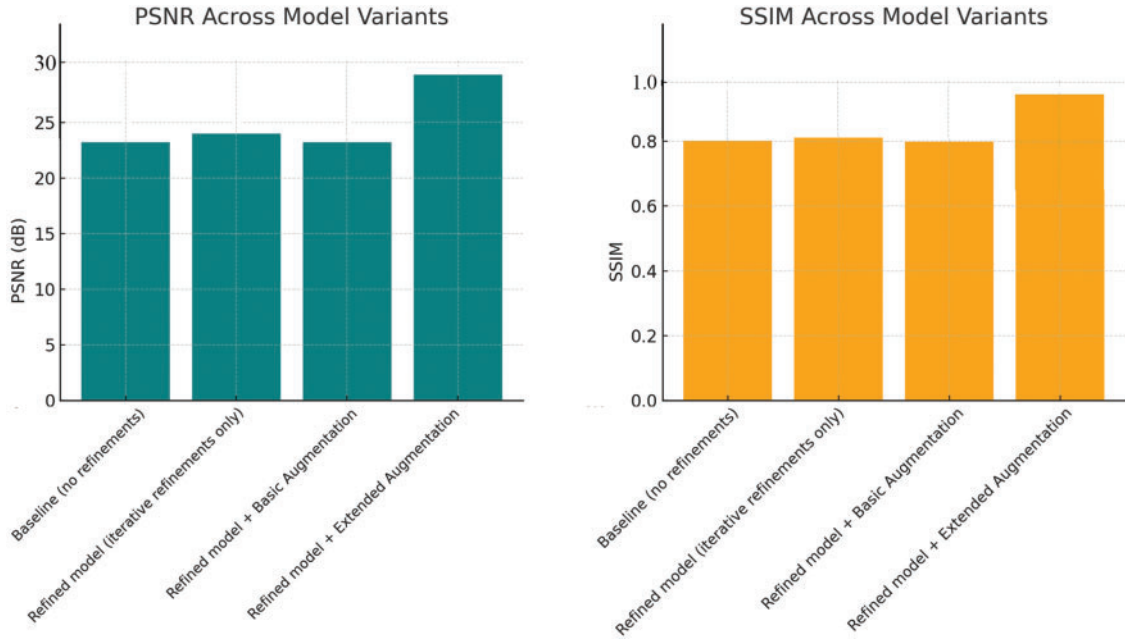


Fig. 5. Comparison of PSNR and SSIM across model variants.

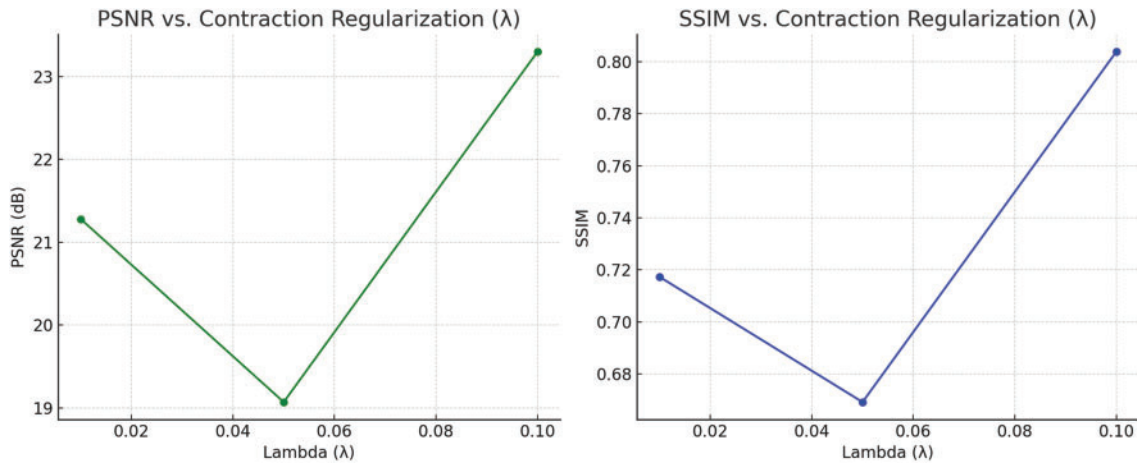


Fig. 6. PSNR (left) and SSIM (right) vs. contraction regularization.

flexibility but lacked sufficient regularization to stabilize the outputs, achieving a modest PSNR of 21.28 dB and SSIM of 0.7173. These findings confirm that moderate contraction regularization strikes an optimal balance between noise suppression and the preservation of visual detail. These trends are illustrated in Fig. 6.

4.6. Impact of Noise Intensity and Distribution on Image Reconstruction Performance

Fig. 7 illustrates how the model performance, measured by PSNR and SSIM, is affected by varying degrees of Gaussian and burst noise. For Gaussian noise, as the noise factor increases from 0.05 to 0.20, both PSNR and SSIM steadily decrease, indicating that higher noise significantly affects the image reconstruction quality, reducing both pixel-level accuracy and perceptual similarity.

A similar trend was observed for burst noise: as the block size of the noise increased, both the PSNR and SSIM values sharply declined. This confirms that larger burst errors degrade image reconstruction more severely. These findings emphasize the sensitivity of the model to

different noise intensities and distributions, underscoring the importance of training with diverse and challenging noise scenarios to enhance the robustness in real-world wireless image transmission.

For Gaussian noise, increasing the noise factor leads to gradual declines in both PSNR and SSIM. For burst noise, increasing block size results in sharper drops, with a more severe impact on SSIM. This indicates that burst noise introduces stronger perceptual degradation compared to Gaussian noise.

4.7. Qualitative Evaluation and Discussion

This study qualitatively assessed the reconstructed images and their error maps to understand the practical effectiveness of our proposed model against the baseline approach. The visual outputs indicated noticeable improvements with the enhanced model, particularly owing to the inclusion of dynamic contraction regularization guided by Banach space principles. Error maps provided a clear illustration of these enhancements, showing significant reductions in reconstruction errors around

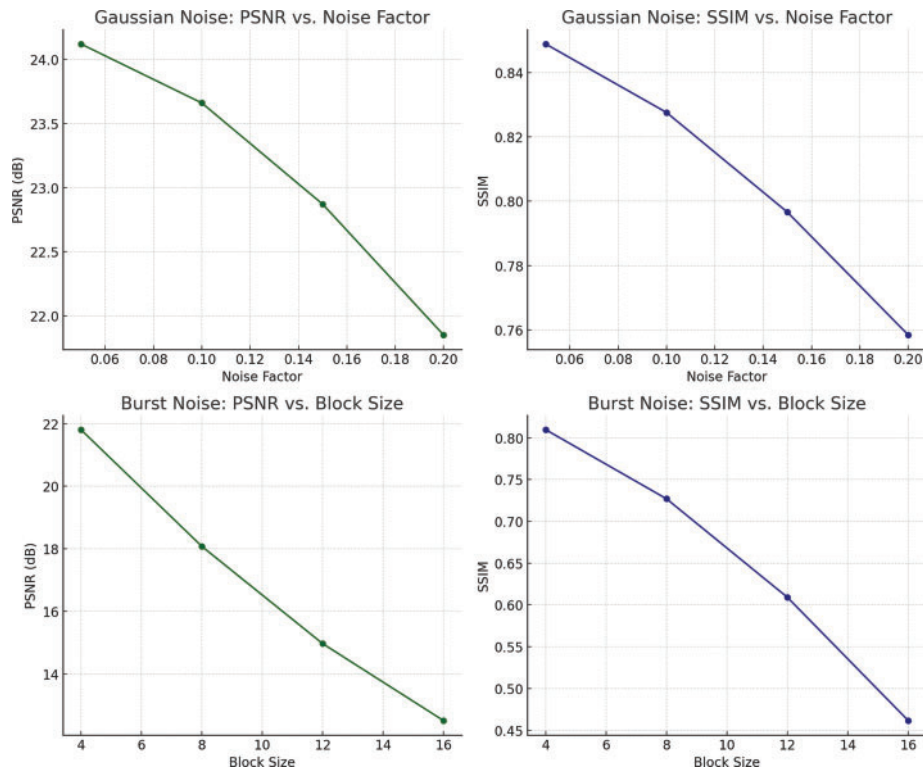


Fig. 7. Figure: Model performance under varying noise conditions.

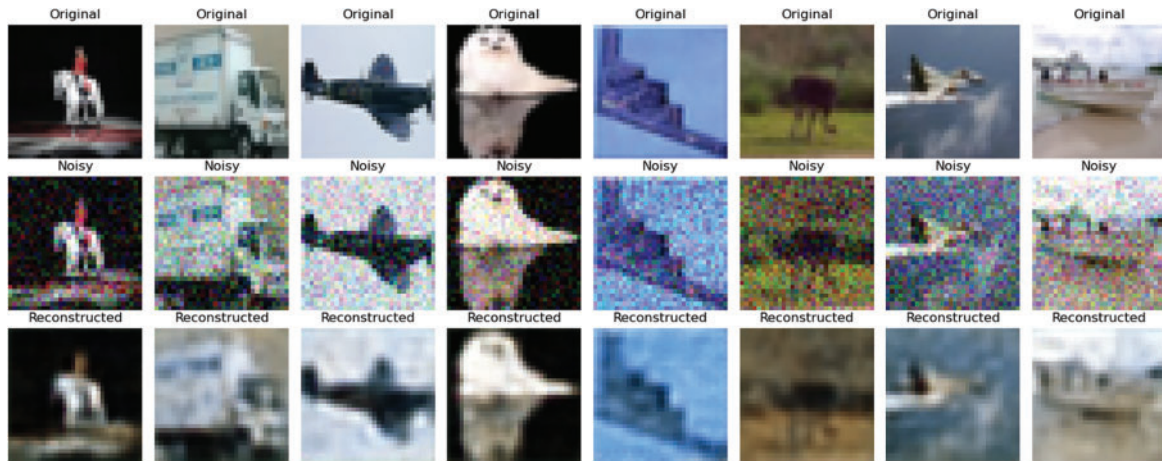


Fig. 8. Comparative reconstruction scenario using the baseline autoencoder.

critical features, such as edges and textured regions. The images produced by the refined model appear sharper, clearer, and structurally more accurate. The qualitative superiority of our approach is particularly pronounced under challenging conditions involving Gaussian and burst noise.

In contrast, the baseline model outputs frequently exhibited blurred features and structural distortions, particularly under higher noise intensities. The image reconstructions presented in this paper provide visual comparisons that distinctly illustrate these improvements, showcasing how the enhanced model consistently delivers better perceptual results. These qualitative gains directly correlated with the quantitative SSIM improvements noted earlier, confirming the model's enhanced alignment with human visual perception.

Specifically, our analysis provides several critical insights. First, the iterative refinement approach significantly enhances the reconstruction quality by progressively correcting distortions and driving the network toward convergence. Second, the implementation of advanced data augmentation techniques that closely mimic real-world noise substantially improves the generalization capability of the model. Third, incorporating moderate contraction regularization grounded in Banach space theory improves the robustness of the model, prevents overfitting, and ensures stable, consistent outputs. Finally, dynamically adjusting the strength of contraction regularization throughout the training effectively balances the accuracy with noise resilience.

Overall, integrating theoretical concepts from functional analysis into practical deep learning strategies has proven to be highly effective for achieving robust, high-quality

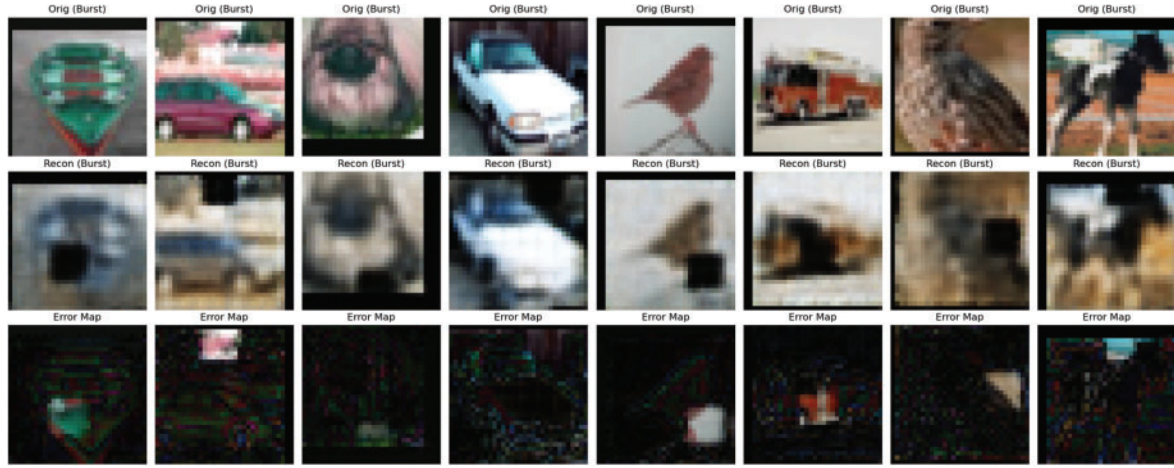


Fig. 9. Burst noise reconstruction and error maps.

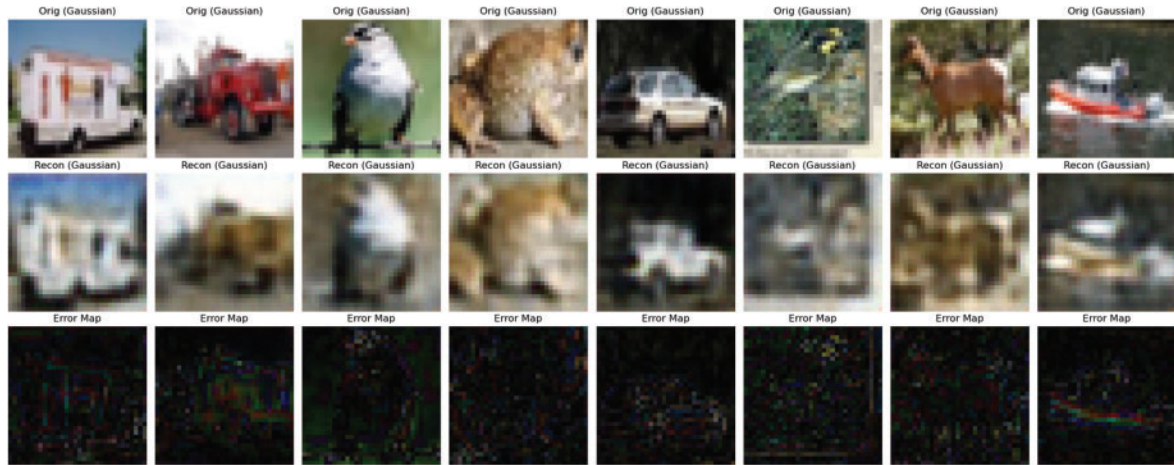


Fig. 10. Gaussian noise reconstruction and error maps (Baseline model).

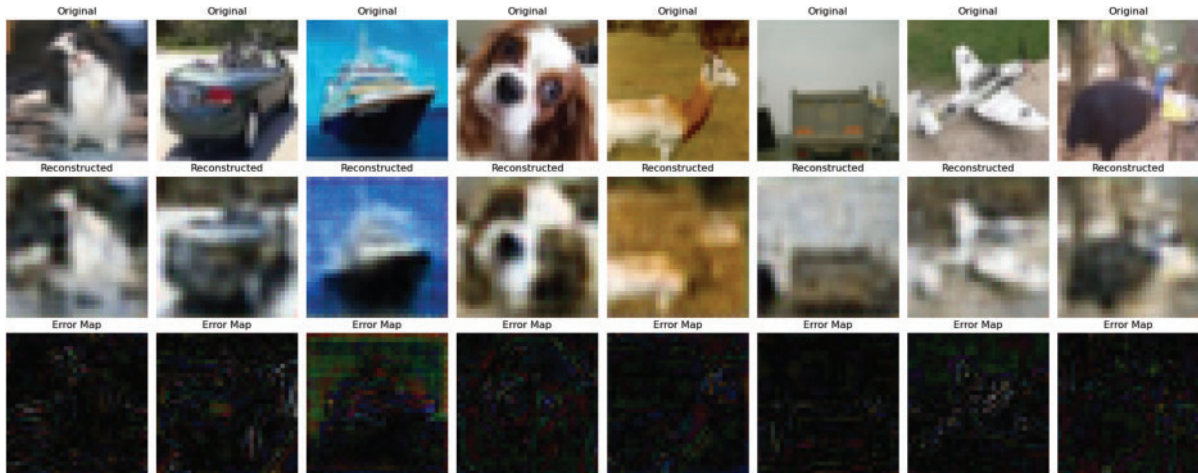


Fig. 11. Reconstruction performance for CIFAR-10 images where λ set at 0.10.

real-time image transmission. These findings not only validate our methodological approach but also suggest a promising direction for future research and practical applications in environments prone to transmission noise.

Fig. 8 illustrates a comparative reconstruction scenario using the baseline autoencoder. The original CIFAR-10 images, reconstructed versions, and associated error maps were presented. The reconstruction quality was noticeably inferior compared to the models enhanced by Banach

space integration and iterative refinement. Error maps reveal widespread distortions, emphasizing the limited capability of the baseline model to handle complex noise conditions, thus underscoring the need for enhanced methodologies.

Fig. 9 shows CIFAR-10 images corrupted by burst noise along with reconstructions generated by the autoencoder enhanced with dynamic contraction regularization and advanced augmentation. The bottom row contains error

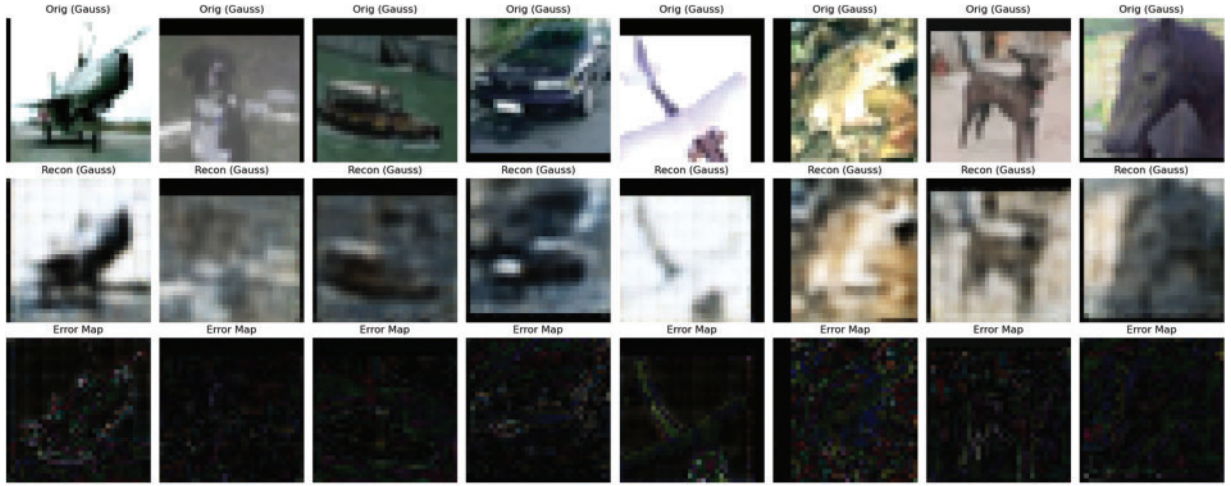


Fig. 12. Gaussian noise reconstruction and error maps (Dynamic contraction with advanced augmentation).

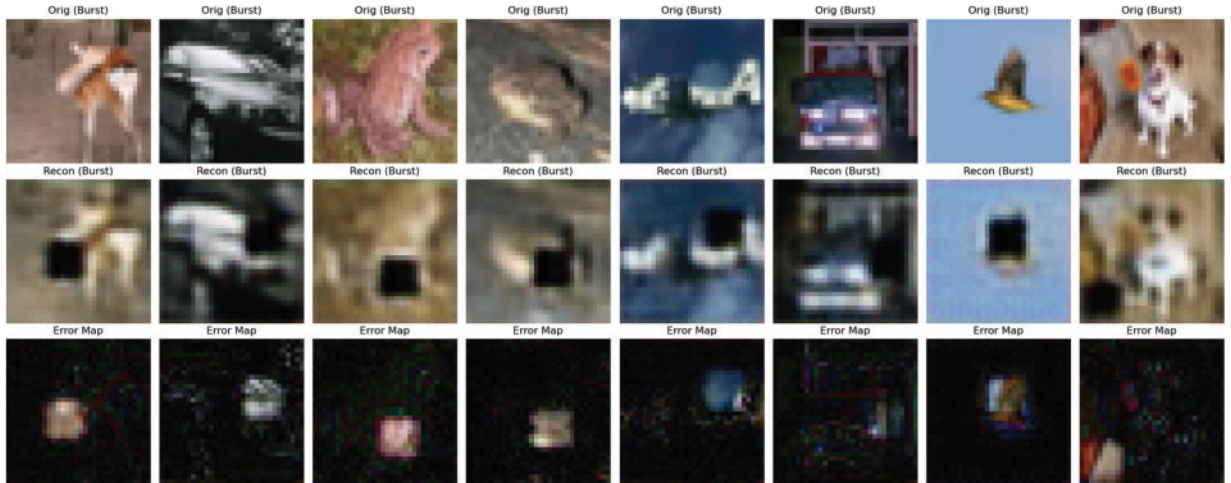


Fig. 13. Visual qualitative analysis of the reconstruction performance of the autoencoder under burst noise conditions.

maps illustrating the residual reconstruction errors. The enhanced model effectively restores the images, significantly reducing the visual distortions caused by burst noise. The error maps highlight the remaining minor errors, primarily around the edges and the detailed textures. This underscores the effectiveness of iterative refinement.

Fig. 10 compares the original images corrupted with Gaussian noise and their reconstructions generated by the baseline autoencoder. The error maps provided here clearly reveal significant residual noise and reconstruction errors, which are particularly noticeable around the key image details. These results reinforce the advantages of integrating the dynamic contraction and iterative refinement techniques into the model architecture.

Fig. 11 shows the reconstruction performance for CIFAR-10 images using a contraction-regularized model with λ set to 0.10. The reconstructions presented notable visual quality improvements relative to the baseline models, effectively mitigating noise-induced artifacts. The error maps highlight the improved precision around image structures, clearly demonstrating the benefits of a moderate regularization factor in balancing noise suppression and detail preservation.

Fig. 12 shows the original CIFAR-10 images with Gaussian noise, their reconstructions using the refined autoencoder (dynamic contraction and advanced augmentation), and the corresponding error maps. These reconstructions demonstrate significant improvements in image clarity and quality, effectively removing Gaussian noise artifacts. The error maps indicate precise areas of minor inaccuracies, revealing the model's strong performance in accurately restoring subtle image features.

Fig. 13 shows a qualitative visual analysis of the reconstruction performance of the autoencoder under burst noise conditions. The top row displays the original images from the CIFAR-10 dataset affected by burst noise, whereas the middle row shows the corresponding reconstructions generated by the proposed deep convolutional autoencoder. The bottom row presents the associated error maps, highlighting the areas with significant reconstruction errors. The visual analysis clearly demonstrates that the model effectively corrects large blocks of burst noise, significantly restoring the image details. However, challenging regions, particularly around object edges and textured areas, retain some artifacts. These error maps provide critical insight into the model performance, emphasizing the efficacy of iterative refinement and Banach space-based contraction regularization

in substantially reducing visual distortions in noisy wireless image transmissions.

All quantitative metrics, PSNR, SSIM, and inference times, along with the scripts used to generate noise patterns, training logs, and model checkpoints, were hosted at https://github.com/cgkinyua/deep_learning_project. This open archive allows readers to verify our comparisons between baseline and Banach-integrated models and to reproduce every figure and table presented in this paper.

5. CONCLUSION

The primary objective of this study is to develop a convolutional autoencoder model that leverages contraction regularization and iterative refinement to maintain high-quality reconstructions, even under challenging noise conditions. The authors demonstrated the successful integration of Banach space principles into an AI-driven error correction for real-time wireless image transmission. By combining deep learning frameworks with functional analytic concepts, we have achieved significant advancements in both the accuracy and robustness of image reconstruction.

The results show that the proposed model consistently outperformed the baseline autoencoder, demonstrating substantial improvements in both the Peak Signal-to-Noise Ratio (PSNR) and Structural Similarity Index Measure (SSIM). Furthermore, the model maintained an inference speed of 614.47 frames per second, making it highly suitable for real-time applications. The integration of contraction mappings, guided by the Banach Fixed-Point Theorem, proved instrumental in enhancing model stability and ensuring convergence. The iterative refinement process further contributes to error reduction and robust reconstruction, significantly outperforming traditional approaches.

The dynamic contraction regularization strategy, which varied the regularization strength during training, effectively balanced the reconstruction accuracy with the model stability. This combination of methods ensured that the model could maintain a high performance even in the presence of diverse noise scenarios, including Gaussian and burst noise. The authors demonstrated the practical feasibility of the proposed model in a GPU-based setup using an NVIDIA A30 GPU. The real-time processing capability of the model, combined with its enhanced reconstruction quality, positions it as a viable solution for applications such as telemedicine, interactive video streaming, and surveillance systems where reliable image transmission is crucial.

Nevertheless, some limitations of this study persist. The performance of the model was primarily evaluated using the CIFAR-10 dataset, which consists of low-resolution images. Future research will focus on adapting this architecture to handle higher-resolution data more efficiently. Moreover, while simulation-based evaluation provides a controlled environment to assess model robustness, real-world testing under varying transmission conditions remains essential for fully validating the model's practical utility. Additionally, there is the potential to explore more advanced network architectures and hybrid approaches

that could further enhance the performance without compromising real-time processing.

The promising results achieved through this research underscore the potential of combining deep learning with Banach space principles to address error correction challenges in wireless image transmission. Embedding theoretical rigour into practical applications. This study bridges the gap between mathematical theory and real-world implementations. Further advancements could include refining the contraction scheduling mechanism to make it adaptive based on noise level variations, and integrating transformer-based architectures to enhance feature extraction.

This study lays a strong foundation for the future exploration of robust, real-time image transmission solutions driven by cutting-edge mathematical principles and AI innovations. We have made all the research materials for this study public to advance the research. By releasing our full experimental pipeline at https://github.com/cgkinyua/deep_learning_project, we have ensured transparency and facilitated further innovation. We encourage the research community to inspect, reproduce, and build upon our work to advance robust real-time wireless image transmission.

6. RECOMMENDATIONS

This study successfully integrated Banach space principles into deep learning architectures to enhance the error correction for real-time wireless image transmission. Despite achieving promising results in terms of both accuracy and inference speed, some limitations remain, which suggest clear directions for future research. Primarily, the use of the CIFAR-10 dataset, which consists of low-resolution images, limits the scalability of the model to more complex or high-resolution data. Future efforts should investigate how the model performs on large-scale image or video datasets, potentially through the use of multiscale architectures or transformer-based networks to improve feature extraction.

Although simulation-based experiments with Gaussian and burst noise demonstrated the robustness of the model, these scenarios do not fully reflect real-world wireless environments. Therefore, conducting real-time experiments under dynamic, noisy transmission conditions is essential to validate the model's practical applicability in fields such as telemedicine, surveillance, and interactive media. Another area that requires attention is the model's use of manually tuned dynamic contraction regularization. Although effective, its reliance on static scheduling may hinder flexibility. Adopting adaptive regularization techniques that adjust λ in response to real-time noise levels can improve robustness and generalizability during deployment.

Furthermore, although the model ran efficiently on an NVIDIA A30 GPU, deployment on edge devices with limited processing capacity remains a challenge. Research on lightweight architectures and model optimization strategies, such as pruning, quantization, and knowledge distillation, could enable broader deployment without sacrificing accuracy. In conclusion, this study provides a strong foundation for robust and theoretically

grounded wireless image transmissions. Future research can further enhance the practical and technical value of these methods by scaling to high-resolution data, conducting real-world testing, developing adaptive regularization strategies, improving model efficiency, and exploring hybrid architectures.

DECLARATION OF GENERATIVE AI AND AI-ASSISTED TECHNOLOGIES IN THE WRITING PROCESS

During the preparation of this work, the author(s) used ChatGPT for grammar and sentence restructuring. After using this tool, the author(s) reviewed and edited the content as required and took (s) full responsibility for the content of the publication.

DECLARATION OF COMPETING INTEREST

The authors, Charles Kinyua Gitonga and team, declare that no conflicts of interest related to this manuscript.

DATA AVAILABILITY

The data underlying this article are available in the Zenodo repository at <https://doi.org/10.5281/zenodo.15381896>, and the codebase is accessible via GitHub at https://github.com/cgkinyua/deep_learning_project/tree/main/cifar-10-batches-py.

CONFLICT OF INTEREST

The authors declare that they do not have any conflict of interest.

REFERENCES

- [1] Hanzo L. Wireless channel impairments and mitigation strategies. *IEEE Commun Surv Tutorials*. 2022;24(3):121–37.
- [2] Chen KC. *Artificial Intelligence in Wireless Robotics*. Boca Raton, FL, USA: CRC Press; 2022.
- [3] Vahidi M, Wu H. Classical error-correcting codes for wireless transmission. *IEEE Trans Commun*. 2020;68(5):1028–37.
- [4] Naseri M, Ashtari P, Seif M, Poorter EDe, Poor HV, Shahid A. Deep learning-based image compression for wireless communications: impacts on reliability, throughput, and latency. arXiv preprint arXiv:2411.10650, Nov. 2024. Available from: <https://arxiv.org/abs/2411.10650>.
- [5] Zhang X, Cao B, Zhang Y. Deep learning for image restoration in real-time wireless transmission. *IEEE Trans Image Process*. 2019;28(7):1505–16.
- [6] Brézis H. *Functional Analysis: Banach Spaces and Contraction Principles*. Springer; 2010.
- [7] NVIDIA Corporation. *NVIDIA A30 GPU Architecture*. NVIDIA Technical Documentation; 2022.
- [8] Meulen A. Evolution of mobile communications and multimedia transmission. *IEEE Wirel Commun*. 2018;25(1):45–52.
- [9] Lee J, Hong S. Orthogonal frequency division multiple access in 4G systems. *IEEE Trans Commun*. 2020;69(4):1234–45.
- [10] Vahidi M, Wu H. Error-correcting codes for real-time wireless transmission. *IEEE Trans Signal Process*. 2020;68(5):1028–37.
- [11] Gallager R. Low-density parity-check codes. *IRE Trans Inf Theory*. 1962;8(1):21–8.
- [12] Kurose J, Ross K. *Computer Networking: A Top-Down Approach*. 8th ed. Pearson; 2021.
- [13] Boursoulatz D, Barbier D, Katsaggelos A. Joint source-channel coding with neural networks. *IEEE J Sel Areas Commun*. 2019;37(6):1291–303.

- [14] Xiao C, Wei W, Liu M. Learning-based wireless image transmission: adapting to noisy channels. *IEEE Commun Lett*. 2021;25(2):510–4.
- [15] Brézis H. *Functional Analysis, Sobolev Spaces and Partial Differential Equations*. Springer; 2010.
- [16] Wang Z, Bovik AC, Sheikh HR, Simoncelli EP. Image quality assessment: from error visibility to structural similarity. *IEEE Trans Image Process*. 2004 Apr;13(4):600–12.
- [17] Hore A, Ziou D. Image quality metrics: PSNR vs. SSIM. *Proc. 20th Int. Conf. Pattern Recognition (ICPR)*, pp. 2366–9, Istanbul, Turkey, 2010.
- [18] Zhang R, Isola P, Efros AA, Shechtman E, Wang O. The unreasonable effectiveness of deep features as a perceptual metric. *Proc. IEEE Conf. Computer Vision and Pattern Recognition (CVPR)*, pp. 586–95, 2018.
- [19] ITU-T Recommendation P.800. *Methods for Subjective Determination of Transmission Quality*. Geneva, Switzerland: Int. Telecommunication Union; 1996 Aug.
- [20] Esmailzadeh H, Blem E, Amant RSt, Sankaralingam K, Burger D. Dark silicon and the end of multicore scaling. *Proc. 38th Annu. Int. Symp. Computer Architecture (ISCA)*, pp. 365–76, San Jose, CA, USA, 2011.
- [21] Choi Y, El-Khamy M, Lee J. Towards the limit of network quantization. *Proc. Int. Conf. Learning Representations (ICLR)*, 2017.