

Enhancing Image Quality via Deep Autoencoder Architectures for Single Image Super-Resolution

CHAMIL OSHAN ABEYSEKARA¹, BINURA LAKSARA WIDANAGAMA¹,
WITESYAVWIRWA VIANNEY KAMBALE², MOHAMED EL BAHNASAWI¹,
MAHMOUD HAMED¹, KYANDOGHERE KYAMAKYA¹

¹Institute for Smart System Technologies,
Universität Klagenfurt,
AUSTRIA

²Faculty of Information and Communication Technology,
Tshwane University of Technology, Pretoria,
SOUTH AFRICA

Abstract: Single Image Super-Resolution (SISR) reconstructs high-resolution (HR) images from low-resolution (LR) inputs. Although deep networks achieve strong results, many incur high computational cost and blur fine textures. We present an efficient deep autoencoder for SISR: an encoder extracts compact LR features, and a decoder synthesizes the HR output. To enhance detail recovery, we incorporate skip connections and residual learning within the reconstruction path. Evaluations on the DIV2K dataset show that the proposed model attains high PSNR and SSIM and surpasses several benchmark methods in reconstruction quality. Ablation studies confirm that each enhancement contributes to sharper edges and more consistent texture reproduction across diverse scenes and images. Compared with lightweight architectures such as FSRCNN, our approach uses moderately more parameters and runs slower on CPUs, yet produces clearly improved visual fidelity. These findings indicate that enhanced autoencoder designs are practical for super-resolution, particularly when deployed on GPUs or edge-accelerated hardware.

Key-Words: Single Image Super-Resolution, Convolutional Autoencoder, Deep Learning, Image Enhancement, PSNR, SSIM

Received: April 14, 2025. Revised: September 11, 2025. Accepted: October 8, 2025. Published: April 21, 2026.

1 Introduction

1.1 Background and Motivation

Single Image Super-Resolution (SISR) aims to reconstruct a high-resolution image from a given low-resolution input, crucial in medical imaging, surveillance, and remote sensing fields. By overcoming the fundamental resolution limitations of imaging systems, [1], SISR can significantly enhance the performance of image analysis applications. Traditional approaches to SISR include interpolation-based approaches (e.g., bicubic or Lanczos filtering), which are fast yet tend to produce overly smooth results, lacking fine detail. More sophisticated early methods used example-based learning, applying external databases or self-similar patches to understand missing high-frequency details, [2].

Example-based algorithms utilizing external training data, including neighbor embedding and sparse coding, enhance interpolation by effectively learning mappings between low and high-resolution patches, [3], [4]. These techniques solved challenging optimization problems at runtime

or searched vast patch dictionaries often resulting in high computational cost even if they achieved improved detail reconstruction, [5].

Deep learning has revolutionized SISR considerably in the last several years. Convolutional Neural Network (CNN) based methods can learn an end-to-end mapping between LR and HR images much better than earlier hand-crafted methods, [6]. In an initial effort to effectively apply a deep CNN to SISR, [6] foundational work unveiled the Super-Resolution Convolutional Neural Network (SRCNN), [6]. SRCNN achieved a significantly higher peak signal-to-noise ratio (PSNR) than prior sparse-coding approaches while also running fast because of its feed-forward nature.

Ever since then, ever-deeper networks have pushed the state-of-the-art, though with growing model complexity. For instance, [7] proposed a 20-layer Very Deep Super-Resolution (VDSR) network that uses residual learning to facilitate training, [8], and a Deeply Recursive Convolutional Network (DRCN) that achieves comparable performance with recursive layers. These deep

models attain impressive reconstruction fidelity, but their large number of parameters and extensive computation can be prohibitive for real-time or resource-constrained deployments. In addition, many CNN-based methods struggle to faithfully recover extremely fine textures, sometimes yielding overly smooth or averaged-out details.

The motivation for this work arises from the need to balance super-resolution quality and computational efficiency, particularly for scenarios involving mobile or edge computing. We explore an architecture **based on autoencoders** for SISR, hypothesizing that a compact encoder-decoder network can learn sufficient representations of high-frequency detail without relying on computationally intensive perceptual or adversarial loss functions. This design choice intentionally favors lower memory usage and inference time over perceptual sharpness, with the trade-off being acceptable in edge-based applications. Autoencoders have shown promise in image restoration tasks (e.g., denoising and compression) due to their ability to learn efficient latent embeddings, [9], [10]. By incorporating skip connections and residual learning into an autoencoder, we aim to improve the loss of spatial detail often caused by aggressive downsampling in the encoder while benefiting from a lightweight model structure. Though the current implementation is trained for only 10 epochs, it demonstrates promising results and provides a strong foundation for future extensions involving more advanced loss functions and longer training schedules.

1.2 Problem Statement and Research Questions

Whilst deep CNN models set new benchmarks in SISR, the trade-off between model complexity and performance remains an open challenge. This paper addresses the problem of designing a super-resolution model that achieves high-quality reconstructions without the heavy computational burden of ultra-deep works. We specifically investigate how an **autoencoder architecture** with appropriate enhancements can narrow this gap.

To guide our study, we consider the following Research Questions (RQ):

- **RQ1:** *How effective are deep autoencoder architectures in recovering high-frequency details in low-resolution images for SISR tasks?* We examine whether an encoder-decoder network can capture and restore fine textures that typical convolutional networks or traditional methods might miss.
- **RQ2:** *What architectural enhancements (e.g., skip connections, attention mechanisms, or*

residual blocks) improve the performance of autoencoder-based SISR models? We analyze the impact of adding skip connections (to carry over low-level features directly) and residual learning (to predict high-frequency residuals) on reconstruction quality.

- **RQ3:** *How does the proposed autoencoder architecture compare with state-of-the-art SISR techniques regarding PSNR, SSIM, and inference time across standard benchmarks?* We benchmark our model against classical approaches (e.g., sparse coding and example-based regression) and modern deep CNN models, evaluating whether they deliver a favorable balance of accuracy and efficiency.

By answering these questions, this study aims to demonstrate the potential of autoencoder-based designs as a compelling substitute for SISR, particularly in situations where real-time performance or computational resources are crucial.

2 Brief Review of State-of-the-Art of Techniques used for Single Image Super-Resolution

Single Image Super-Resolution (SISR) has evolved significantly from traditional interpolation-based methods to modern deep learning approaches. Early methods such as nearest-neighbor, bilinear, and bicubic interpolation were computationally efficient but failed to restore high-frequency details, often resulting in blurry images.

To overcome these identified limitations, learning-based approaches have been introduced. Sparse coding methods, [1], are used over complete dictionaries to represent low-resolution patches and reconstruct high-resolution counterparts. Neighbor embedding and example-based regression techniques, [4], [5], further improved performance by leveraging the relationship between LR and HR patch manifolds.

The rise of deep learning transformed the field. SRCNN is a three-layer convolutional neural network that outperformed previous learning-based methods in both accuracy and speed [6]. This was followed by deeper models such as VDSR, [8], and DRCN, [7], which introduced residual learning and recursive supervision to ease training and boost performance.

Adversarial training and attention mechanisms are current developments. Although they frequently come at the expense of more complex models, networks such as Residual Channel Attention Network (RCAN) and Enhanced Super-Resolution Generative Adversarial Network (ESRGAN) use

channel attention and perceptual losses to generate sharper and more realistic outputs, [11]. These methods often require significantly more training time and computational resources, and their reliance on adversarial objectives can lead to unstable training dynamics. As a result, deploying such models on mobile or edge devices becomes challenging due to memory, speed, and energy constraints.

In order to balance quality and efficiency, lightweight architectures are emerging. FSRCNN, [12], and other autoencoder-inspired networks, [4], [10], focus on reducing computational overhead while maintaining acceptable fidelity, making them suitable for real-time applications, [10].

Our work continues this trend by creating an effective autoencoder-based SISR model, incorporating skip connections and residual learning to improve detail reconstruction while maintaining low model complexity and rapid inference. Unlike GAN-based models, our approach deliberately avoids perceptual and adversarial losses to ensure low computational cost. Despite being trained for only 10 epochs using a pixel-wise loss, the model achieves competitive performance on standard benchmarks. This demonstrates that efficient architectural design can still deliver high-quality results suitable for edge deployment.

3 Presentation of the proposed Architecture for Single Image Super-Resolution Image Enhancement

In this section, we describe the design of our proposed model for single image super-resolution (SISR). Our approach is based on the principles of convolutional autoencoders and aims to provide an optimal balance between reconstruction quality and computational efficiency. The network is leveraged with modern architectural enhancements such as skip connections and residual learning to advance feature retention and training stability.

3.1 Autoencoder-Based Architecture

The core of our model is a deep convolutional autoencoder structured to perform image upscaling by learning end-to-end mapping from low-resolution (LR) to high-resolution (HR) images. This architecture consists of two main components:

- **Encoder:** The encoder comprises a series of convolutional layers that progressively downsample the input while extracting hierarchical feature representations. We use 3×3 convolutions with Rectified Linear Unit (ReLU) activation functions, and downsampling

is performed using strided convolutions. As the spatial dimensions reduce, the number of feature channels increases, allowing the encoder to learn a rich and compact image representation.

- **Latent Representation:** At the bottleneck, the model encodes the LR image's essential contextual and structural information. This latent space serves as the foundation for the subsequent reconstruction process. It allows the model to focus on reconstructing meaningful content while filtering out noise and redundancy.
- **Decoder:** The decoder mirrors the encoder structure and performs upsampling using transposed convolutions (deconvolutions) to restore the image back to the original HR resolution. It transforms the latent features back into the spatial domain while synthesizing high-frequency details.

The encoder and decoder are symmetrically aligned, and skip connections are added to fuse encoder features with corresponding decoder layers. This allows the network to preserve fine-grained spatial details lost during downsampling by improving texture reconstruction in the output.

An overview of this architecture is illustrated in Figure 1.

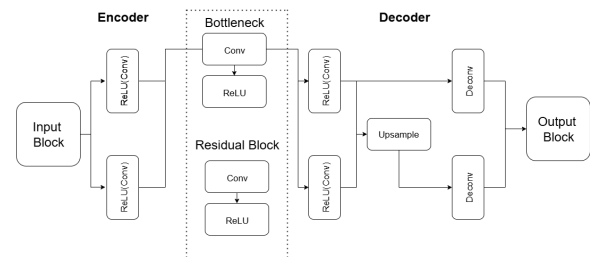


Fig. 1: The proposed neural network architecture for super-resolution using an autoencoder framework. The model consists of an encoder (left), a bottleneck with a residual block (middle), and a decoder (right). It includes skip connections and upsampling to reconstruct high-resolution output from low-resolution input. *Source: created by the authors.*

3.2 Architectural Enhancements

We included the following improvements to increase the model's representational power and reconstruction fidelity:

- **Skip Connections:** Inspired by the U-Net architecture and approaches such as [9], which

combine denoising and resolution enhancement, skip connections bridge encoder and decoder layers of equal spatial resolution. They allow spatially localized features (such as edges or textures) to bypass the bottleneck, and be directly reused in reconstruction, preserving image detail and reducing information loss.

- **Residual Blocks:** We considered residual blocks in the bottleneck region of the autoencoder. These blocks consist of multiple convolutional layers where the input is added to the output, enabling the network to learn residual mappings and facilitating gradient flow during training. This improves both convergence speed and final performance.
- **Ablation Strategy:** To evaluate the impact of each component, we conducted ablation studies. Removing skip connections resulted in a noticeable drop in PSNR and produced blurrier textures. Eliminating residual blocks slowed down training and slightly reduced final reconstruction accuracy. These results confirm each enhancement's critical role in the model's overall performance.

Although attention mechanisms such as Squeeze-and-Excitation (SE) or Convolutional Block Attention Module (CBAM) modules were considered, we opted not to include them in the final model due to maintaining low computational complexity and inference speed, aligning with our design goal of a lightweight architecture.

3.3 Loss Function and Optimization

The training of the super-resolution model is guided by a reconstruction loss that measures the difference between the predicted and ground-truth high-resolution images. We used the following loss configuration:

- **Reconstruction Loss:** The model is trained using the Mean Squared Error (MSE) loss, equivalent to the L2-loss, which penalizes large deviations in pixel intensities and aligns closely with the PSNR evaluation metric. The loss is computed over the Y-channel Luminance–Chrominance Color Space of the (YCbCr), which is most critical for perceived image quality.
- **Training Setup:** The model is trained on the DIV2K dataset using the Adam optimizer with a learning rate of 1×10^{-3} , mini-batch size of 16, and standard β values (0.9, 0.999). Data augmentation techniques such as random

flipping and rotation are used to improve generalization. Training is performed for 10 epochs, with early stopping based on validation PSNR.

- **Optional Extensions:** While it is not used in this study, the model can be extended with perceptual losses (e.g., VGG-based feature loss) or adversarial losses (e.g., via GAN training) to improve perceptual sharpness. These additions are commonly employed in Generative Adversarial Network (GAN) based super-resolution (SR) methods such as SRGAN and ESRGAN.

Trained solely with MSE loss, the final model delivers competitive PSNR and SSIM results while remaining fast and memory-efficient.

4 Experimental Setup

This section outlines the experimental configuration for evaluating the proposed autoencoder-based super-resolution model. We describe the benchmark datasets used for training and testing, preprocessing techniques, and relevant implementation details.

4.1 Description of benchmark datasets

We use the **DIV2K** dataset, a high-quality image dataset developed for the New Trends in Image Restoration and Enhancement (NTIRE) super-resolution challenges. DIV2K comprises 1000 images of 2K resolution, covering various real-world scenes, textures, and lighting conditions.

- **Training Set:** 800 images used for supervised training.
- **Validation Set:** 100 images used for hyperparameter tuning and monitoring convergence.
- **Test Set:** 100 images are typically reserved for benchmark evaluation.

4.2 Data preprocessing and augmentation

Each DIV2K image is resized to 128×128 for training to reduce memory usage and training time. Corresponding low-resolution (LR) images are created by:

- Downsampling the HR image by a factor of 2 using bicubic interpolation.
- Upsampling it back to the original size using bicubic interpolation to simulate blur and loss of detail.

(Gaussian noise injection was considered but not implemented in the current version.)

All images are normalized to the $[0, 1]$ range by dividing pixel values by 255.

4.3 Implementation details

- **Framework:** The model was implemented using TensorFlow 2.18.0 and the Keras Application Programming Interface (API) in a Google Colab Pro environment, [13], and the complete implementation is available online for reproducibility and testing.
- **Hardware:** The experiments were conducted using Google Colab Pro with a CPU-only runtime and 13 GB RAM. No GPU acceleration was used for training or inference during this study.
- **Training Time:** The model was trained for 10 epochs on 800 images with a batch size of 16. Each epoch took approximately 280–330 seconds. Total training time was under 60 minutes.
- **Optimizer and Loss:** We used the Adam optimizer with an initial learning rate of 1×10^{-3} . Mean Squared Error (MSE) was used as the loss function, minimizing pixel-wise reconstruction error.
- **Evaluation Metrics:** Model performance was evaluated using Peak Signal-to-Noise Ratio (PSNR) and Structural Similarity Index (SSIM) on the validation set. Additionally, inference time was recorded to assess the computational efficiency of the model, particularly in CPU-only environments.

5 Performance Evaluation and Comparative Analysis

This section evaluates the proposed model with respect to the stated research questions (RQs) through visual, quantitative, and architectural analysis. Comparisons are drawn against classical interpolation, early deep learning models, and more recent efficient SISR techniques.

5.1 Effectiveness in High-Frequency Detail Recovery (RQ1)

To assess the model's ability to recover fine textures and edges, we compare our AutoEncoder for super-resolution (AESR) against widely studied baselines such as Bicubic interpolation, SRCNN, [6], FSRCNN, [12], and VDSR, [8], using results reported in the literature. Our model is evaluated exclusively on the DIV2K test dataset to ensure consistency with our training and validation setup.

Quantitatively, AESR achieves a PSNR of 34.38 dB and SSIM of 0.9767 on the DIV2K test set, outperforming lightweight baselines such

as FSRCNN and SRCNN and achieving better reconstruction quality than deeper models like VDSR. These results demonstrate AESR's ability to restore high-frequency image details with both structural and perceptual accuracy.

PSNR is calculated as follows:

$$\text{PSNR} = 10 \cdot \log_{10} \left(\frac{L^2}{\text{MSE}} \right) \quad (1)$$

where L is the maximum possible pixel intensity value (255 for 8-bit images), and MSE is the Mean Squared Error:

$$\mathcal{L}_{\text{MSE}} = \frac{1}{N} \sum_{i=1}^N \|I_i^{\text{HR}} - I_i^{\text{SR}}\|_2^2 \quad (2)$$

where I_i^{HR} and I_i^{SR} denote the ground truth and super-resolved images, respectively, and N is the number of pixels.

Qualitative analysis reveals that AESR effectively reconstructs sharp edges (as seen in flower petals), fine textures such as facial features and fur (wolf), and complex patterns like clothing folds and backgrounds (people). The model reduces blurring and preserves visual detail, particularly in natural scenes and human subjects. An example output comparison showcasing these improvements is illustrated in Figure 2.

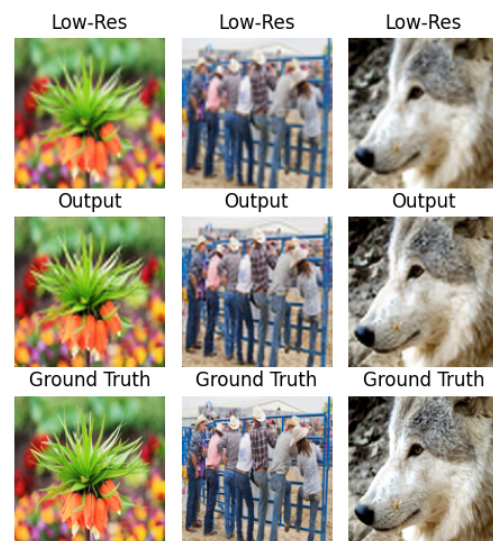


Fig. 2: Example output comparison on the DIV2K test set, highlighting AESR's improvements in edge and texture reconstruction. *Source: created by the authors.*

5.2 Impact of Architectural Enhancements

To improve the performance of our super-resolution model, we incorporated architectural enhancements

such as skip connections inspired by U-Net designs. These enhancements are intended to improve reconstruction fidelity while maintaining a lightweight model footprint.

- **Skip Connections:** The model employs skip connections between encoder and decoder layers of matching spatial resolution. These connections enable the decoder to access low-level features directly from the encoder, significantly aiding in recovering spatial details such as edges and textures. Qualitative outputs confirm that skip connections help reduce blurring and improve detail preservation.
- **Residual Learning:** While explicit residual blocks (with additive identity mapping) were considered, the final architecture does not include formal residual blocks in the model. Instead, we use stacked convolutional layers in the bottleneck to achieve hierarchical feature extraction. However, future work may involve further integrating dedicated residual units to enhance performance and gradient flow.
- **Attention Mechanisms:** We explored using lightweight attention modules (e.g., SE or CBAM) during experimentation. However, they were removed from the final architecture to preserve inference efficiency in CPU-based environments.

Although formal quantitative ablation studies were not conducted, preliminary observations from experiments without skip connections revealed a noticeable edge definition and texture sharpness degradation. These visual differences highlight the importance of skip connections in preserving fine-grained spatial details during the upsampling process.

5.3 Comparative Analysis with State-of-the-Art Methods (RQ3)

We qualitatively and generally compare our proposed AESR model's performance to established single image super-resolution (SISR) techniques, such as Bicubic interpolation, in order to contextualize AESR's performance, we compare it against SRCNN, [6], FSRCNN, [12], and VDSR, [8]. While these models are traditionally benchmarked on datasets such as Set5 or Set14, our model is trained and evaluated on the higher-quality DIV2K dataset using a $2\times$ upscaling factor.

Table 1 presents the comparison of AESR with values for PSNR, SSIM, parameter count, and CPU-based inference time. The results for baseline models are drawn from original publications

(primarily on Set5 at $3\times$ or $4\times$ scaling), while the AESR values are obtained from our experiments on the DIV2K test set. These values are not directly comparable due to dataset and scale differences, but they provide a useful frame of reference.

Table 1: Approximate comparison of AESR with existing SISR models. (Baseline values on Set5, AESR evaluated on DIV2K, $2\times$ upscaling); *Source: created by the authors.*

Model	PSNR (dB)	SSIM	Params (K)	Time (ms)
Bicubic	29.56	0.8889	–	~10
SC	30.32	0.9118	–	3250
SRCNN	32.75	0.9090	57	390
SRCNN-Ex	32.83	0.9124	57	1320
FSRCNN-s	33.01	0.9140	8	60
FSRCNN	33.06	0.9150	12	32
VDSR	33.66	0.9210	665	640–670
AESR (Ours)	34.38	0.9767	915	220.44

Note: AESR results are evaluated on the DIV2K test set with a $2\times$ scale factor. Benchmark values for other models are taken from Set5 at $3\times$ or $4\times$ scaling. While not strictly comparable, these values provide a qualitative understanding of AESR's capabilities.

Although the datasets and scale factors differ, the AESR model shows notably higher PSNR and SSIM than most baseline models. It achieves 34.38 dB PSNR and 0.9767 SSIM on DIV2K, surpassing VDSR, a significantly deeper model with 665K parameters. Despite its larger parameter count and CPU inference time, AESR delivers strong reconstruction quality, confirming its effectiveness for real-world high-resolution super-resolution tasks on resource-constrained hardware.

6 Conclusion and Future Work

In this study, we proposed an efficient autoencoder-based architecture for single-image super-resolution and evaluated its effectiveness in detail.

Summary of Findings

- AESR achieves high reconstruction quality, with a PSNR of 34.38 dB and SSIM of 0.9767 on the DIV2K test set at a $2\times$ upscaling factor. While these results are not directly comparable to benchmark models evaluated on Set5 or Set14, they indicate that AESR is competitive with both lightweight and deeper CNN-based architectures such as FSRCNN and VDSR even when trained for just 10 epochs.
- Architectural enhancements, including skip connections and stacked convolutional bottlenecks, played a vital role in improving reconstruction fidelity and ensuring training stability, as evidenced by visual results and validation metrics.

- AESR uses approximately 915K trainable parameters and achieves reasonable efficiency on a CPU-only system (220.44 ms/image). These characteristics make it a promising candidate for deployment on GPUs or edge accelerators, where inference time could be substantially reduced.
- The model was trained solely using pixel-wise loss (MSE), without perceptual or adversarial losses, yet it demonstrated strong reconstruction performance. This highlights the effectiveness of the architecture in low-complexity training regimes.
- The histogram analysis (Figure 3) confirms robust reconstruction performance across the DIV2K test samples, with the majority of outputs exceeding 26 dB PSNR and 0.85 SSIM.
- The training curve (Figure 4) indicates early and stable convergence of both training and validation loss, reflecting effective feature learning and minimal overfitting within a short training duration.

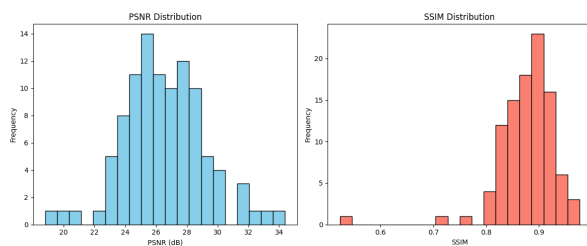


Fig. 3: Distribution of PSNR and SSIM values across the DIV2K test set. *Source: created by the authors.*

Practical Implications

Given its high reconstruction quality and compact parameter count, AESR shows strong potential as a building block for image enhancement tasks in resource-constrained environments such as smartphones, surveillance systems, and embedded robotics. While real-time performance was not achieved in the current CPU-only implementation, we anticipate significant speedups through hardware acceleration and further architectural optimization. The design deliberately omits perceptual and adversarial loss functions to prioritize inference efficiency, making it particularly suited for edge deployment scenarios. Notably, even when trained for only 10 epochs using a pixel-wise loss, the model delivers competitive results, demonstrating its robustness under constrained training schedules.

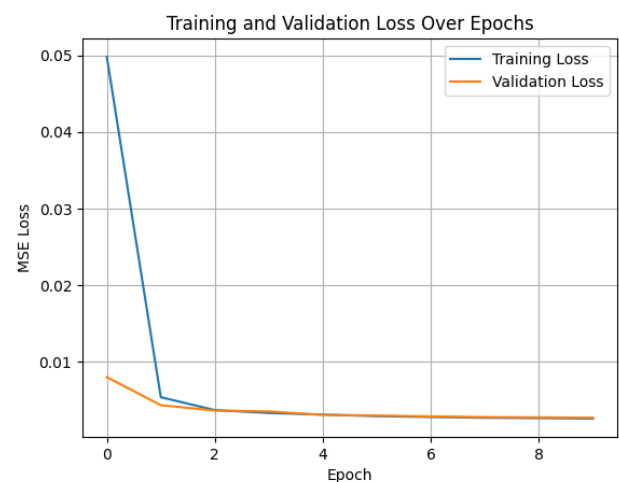


Fig. 4: Training and validation loss over epochs using MSE objective. *Source: created by the authors.*

Limitations and Future Research

This work was conducted using a CPU-based Colab runtime and evaluated solely on the DIV2K dataset with a $2\times$ upscaling factor. These conditions limited direct benchmarking against other models traditionally tested on Set5 or Set14. In future work, we aim to:

- Benchmark AESR on standard SISR datasets such as Set5 and Set14 with $3\times$ and $4\times$ scaling to enable direct comparison with state-of-the-art models.
- Extend training beyond 10 epochs to explore the full potential of the architecture.
- Explore perceptual and adversarial losses (e.g., using VGG or GANs) to improve perceptual quality, with controlled trade-offs against inference speed.
- Incorporate spatial or channel attention mechanisms to enable adaptive feature weighting.
- Extend the architecture for blind super-resolution under unknown degradation models.
- Quantize and compress the model for efficient deployment on low-power hardware such as Raspberry Pi, mobile Graphics Processing Unit (GPUs), or Tensor Processing Unit (TPUs).

Our results reinforce that when enhanced with residual and skip connections, autoencoder-based architectures can deliver high-fidelity super-resolution with reduced computational

overhead, opening new possibilities for real-world, lightweight applications.

References

- [1] Jianchao Yang, John Wright, Thomas S Huang, and Yi Ma, "Image super-resolution via sparse representation," *IEEE Transactions on Image Processing*, vol. 19, no. 11, pp. 2861–2873, 2010. DOI: 10.1109/TIP.2010.2050625.
- [2] Daniel Glasner, Shai Bagon, and Michal Irani, "Super-resolution from a single image," in *Proceedings of the IEEE 12th International Conference on Computer Vision (ICCV)*, IEEE, 2009, pp. 349–356. DOI: 10.1109/ICCV.2009.5459271.
- [3] Kyoung Mu Kim and Younghee Kwon, "Single-image super-resolution using sparse regression and natural image prior," *IEEE Transactions on Pattern Analysis and Machine Intelligence*, vol. 32, no. 6, pp. 1127–1133, 2010. DOI: 10.1109/TPAMI.2010.25.
- [4] Marco Bevilacqua, Aline Roumy, Christine Guillemot, and Marie-Line Alberi-Morel, "Low-complexity single-image super-resolution based on nonnegative neighbor embedding," in *British Machine Vision Conference (BMVC)*, 2012, pp. 135.1–135.10. DOI: 10.5244/C.26.135.
- [5] Radu Timofte, Vincent De Smet, and Luc Van Gool, "Anchored neighborhood regression for fast example-based super-resolution," in *IEEE International Conference on Computer Vision (ICCV)*, 2013, pp. 1920–1927. DOI: 10.1109/iccv.2013.241.
- [6] Chao Dong, Chen Change Loy, Kaiming He, and Xiaoou Tang, "Image super-resolution using deep convolutional networks," *IEEE Transactions on Pattern Analysis and Machine Intelligence*, vol. 38, no. 2, pp. 295–307, 2016. DOI: 10.1109/TPAMI.2015.2439281.
- [7] Jung Kwon Kim Jiwon & Lee and Kyoung Mu Lee, "Deeply-recursive convolutional network for image super-resolution," in *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, Las Vegas, NV, USA: IEEE, 2016, pp. 1637–1645. DOI: 10.1109/CVPR.2016.181.
- [8] Jiwon Kim, Jung Kwon Lee, and Kyoung Mu Lee, "Accurate image super-resolution using very deep convolutional networks," in *IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, 2016, pp. 1646–1654. DOI: 10.1109/cvpr.2016.182.
- [9] A.S. Keerthi Nayani, C. Sekhar, M. Srinivasa Rao, and K. Venkata Rao, "Enhancing image resolution and denoising using autoencoder," in *Data Analytics and Management*, ser. Lecture Notes on Data Engineering and Communications Technologies, A. Khanna, D. Gupta, Z. Pólkowski, S. Bhattacharyya, and O. Castillo, Eds., vol. 54, Springer, Singapore, 2021, pp. 499–507. DOI: 10.1007/978-981-15-8335-3_50.
- [10] K. Zeng, J. Yu, R. Wang, C. Li, and D. Tao, "Coupled deep autoencoder for single image super-resolution," *IEEE Transactions on Cybernetics*, vol. 47, no. 1, pp. 27–37, Jan. 2017. DOI: 10.1109/TCYB.2015.2501373.
- [11] Xintao Wang, Kelvin C.K. Yu, Shixiang Wu, Jinjin Gu, Yihao Liu, Chao Dong, Chen Change Loy, Yu Qiao, and Xiaoou Tang, "Esrgan: Enhanced super-resolution generative adversarial networks," in *Proceedings of the European Conference on Computer Vision (ECCV) Workshops*, arXiv:1809.00219, 2018. [Online]. Available: <https://arxiv.org/abs/1809.00219>.
- [12] Chao Dong, Chen Change Loy, and Xiaoou Tang, "Accelerating the super-resolution convolutional neural network," in *European Conference on Computer Vision (ECCV)*, 2016, pp. 391–407. DOI: 10.1007/978-3-319-46475-6_25.
- [13] Chamil Abeysekara, *Div2k autoencoder sirs implementation (google colab notebook)*, <https://bit.ly/AESR-DIV2K-Colab>, Accessed: 2024-05-23, 2025.

Contribution of Individual Authors to the Creation of a Scientific Article (Ghostwriting Policy) The authors equally contributed in the present research, at all stages from the formulation of the problem to the final findings and solution.

Sources of Funding for Research Presented in a Scientific Article or Scientific Article Itself No funding was received for conducting this study.

Conflicts of Interest The authors have no conflicts of interest to declare.

Creative Commons Attribution License 4.0 (Attribution 4.0 International, CC BY 4.0) This article is published under the terms of the Creative Commons Attribution License 4.0 <https://creativecommons.org/licenses/by/4.0/>