

Lightweight parallel feedback network based on CRL with policy transfer and enhancement for image super-resolution

S V R Manimala, T Kavitha

Department of Electronics and Communication Engineering, MVSR Engineering College, Hyderabad, India

Article Info

Article history:

Received Aug 1, 2025

Revised Apr 6, 2026

Accepted Apr 17, 2026

Keywords:

Curriculum reinforcement learning
EdgeNet
Feedback network
High-resolution image reconstruction
Image super-resolution
Lightweight parallel feedback network

ABSTRACT

Image super-resolution (SR) is essential in applications such as surveillance, medical imaging, and remote sensing, but existing deep learning (DL) models often require high computational resources and struggle to recover fine details in lightweight architectures. Although feedback and attention-based methods have shown improvements, they still lack an effective combination of efficient feature refinement, edge enhancement, and low parameter complexity. To address this gap, we propose a lightweight parallel feedback network (LPFN) that combines three key components: a feedback block for repeated feature refinement, a dispersion-aware attention residual block (DARB) for highlighting important spatial and channel details, and EdgeNet for edge sharpening for sharper boundaries. These components are supported by curriculum reinforcement learning (CRL), an adaptive training strategy that gradually improves the model's learning behavior. Instead of relying on a fixed loss function, LPFN uses a dynamically learned global feedback loss to refine reconstruction quality at each stage. Experiments on DIV2K and Flickr2K show that LPFN achieves higher PSNR and SSIM scores while keeping the model lightweight and efficient. This study emphasizes an effective lightweight feedback framework, an enhanced attention and edge-refinement mechanism, and an adaptive learning strategy that improves both accuracy and stability under different degradation conditions.

This is an open access article under the [CC BY-SA](#) license.



Corresponding Author:

S V R Manimala

Department of Electronics and Communication Engineering, MVSR Engineering College

Hyderabad, Telangana, 501510, India

Email: srikurmam.manimala@gmail.com

1. INTRODUCTION

Image super-resolution (SR) is a pivotal task in visual processing, aimed at reconstructing high-resolution (HR) images from low-resolution (LR) inputs. In encoder-decoder models, the encoder extracts multi-scale features from downsampled LR inputs, while the decoder upsamples and refines these features to generate HR outputs [1]. HR images are vital in areas like surveillance, medical imaging, and remote sensing, where high clarity supports accurate analysis and decision-making [2]. Most SR methods require high computational costs, limiting their application in devices with limited resources [3]. Despite advancements in imaging technologies, capturing and enhancing HR images remains challenging due to: i) the rigidity and cost-prohibitive nature of real-world applications, and ii) the prioritization of capturing new HR images over enhancing existing LR ones.

Early machine learning-based SR methods used handcrafted features and interpolation, which were efficient but failed to recover fine details and lacked scalability. Deep learning (DL), especially convolutional neural networks (CNNs), overcame these limits by learning hierarchical features directly from data,

achieving superior SR performance [4]. Most CNN-based methods emphasize advanced architectures, such as residual learning and dense connections [5]. However, deep networks bring new challenges, including high computational complexity, significant memory requirements, and difficulties in training due to issues like gradient vanishing and explosion as network depth increases [6]. As network depth increases, issues such as gradient vanishing and explosion become prominent, complicating the convergence during training. To enhance resolution performance, researchers have developed deeper networks leveraging residual learning techniques [7]. Traditional DL methods primarily adopt a feed-forward information flow. This approach restricts feedback from later layers to earlier ones, hindering dynamic adjustment of the input at previous layers [8]. To solve these issues, lightweight SR methods have been proposed. However, they often fail to fully exploit the representational capabilities of CNNs.

Recent studies highlight the value of feedback mechanisms for strengthening low-level features using high-level cues in vision tasks [9]. Despite their effectiveness, many SR methods rely on architectures with a large number of parameters, limiting their use in resource-constrained environments. In addition, the need to model long sequence dependencies often result in large numbers of hidden states, causing channel redundancy and hindering the learning of essential SR image representations [10]. This has increased interest in lightweight SR networks [11], where feedback proves especially useful. Recent work on the deep residual channel attention network (RCAN) [12] focuses on optimizing the l_1 loss between the generated HR image and the ground truth. Although it achieves high PSNR, it often produces over-smoothed results with blurred edges. Moreover, studies [13] show that increasing network depth does not always improve performance and can even lead to degradation. To enhance perceptual quality, the super-resolution generative adversarial network (SRGAN) [14] introduced a perceptual loss combining adversarial loss with content loss derived from the VGG network, but this improvement in visual quality comes at the cost of lower PSNR compared to other SR methods. Moreover, existing feedback-based SR methods rely on heavy architectures, and they do not explicitly address the challenge of achieving efficient high-frequency detail recovery in lightweight networks.

While earlier studies have investigated lightweight components, residual learning modules, and feedback mechanisms individually, they have not explicitly examined how these elements behave when combined within a unified lightweight framework especially regarding high-frequency detail recovery and training stability. To address this gap, this study investigates the effects of integrating a lightweight parallel feedback mechanism, dispersion-aware attention, and curriculum reinforcement learning (CRL) within a single SR model. Therefore, this article proposes a novel lightweight parallel feedback network (LPFN) to deliver high-quality SR while maintaining computational efficiency. Initially, the FB is introduced to iteratively refine low-level features using high-level feedback, thereby reducing gradient vanishing and improving convergence. Second, DARB is integrated to selectively focus on informative spatial and channel features, enhancing edge sharpness and fine-texture reconstruction. Then, EdgeNet is employed as a lightweight edge-enhancement module that sharpens contours without adding heavy computational overhead. Finally, a global feedback loss, optimized through CRL with policy transfer and enhancement, is used to dynamically adjust the loss function during training. This adaptive loss learning ensures stable convergence, improved structural consistency, and robustness under varying degradation levels.

The main contribution of this work is detailed in below:

- A novel LPFN is introduced, integrating parallel feedback and residual learning to iteratively refine low-level features using high-level information, enabling high-quality SR with minimal computational overhead.
- An enhanced feature refinement mechanism is developed through DARB and EdgeNet, which selectively emphasize important spatial and channel features and tend to produce a higher proportion of sharp edges and clear texture details in lightweight environments.
- A dynamic learning strategy is incorporated using global feedback loss optimized via CRL, allowing the model to adaptively adjust the learning process, accelerate convergence, and maintain stable reconstruction quality under varying degradation conditions.

2. RELATED WORKS

2.1. Traditional SR methods

Image SR presents a complex inverse problem, as extracting high-frequency details from low LR images is inherently difficult. However, merely increasing the network's depth does not yield significant performance gains and incurs substantial computational costs. To address this, Fang *et al.* [15] developed soft-edge assisted network (SeaNet), which leverages soft-edge features with CNNs to refine SR and produce high-quality images. However, a continuous challenge lies in the complex task of recovering fine texture details in the reconstructed images. To address this challenge, Zhao *et al.* [16] proposed enhanced laplacian pyramid generative adversarial network (ELSRGAN), which uses a Laplacian pyramid to capture

high-frequency details and employs residual-in-residual dense blocks (RRDB) to overcome gradient vanishing. Nonetheless, these approaches may encounter distortions when processing small images that contain significant interference. To address these issues in SR reconstruction, Kong *et al.* [17] proposed enhanced convolutional neural network model (FISRCN), an enhanced CNN for color and texture restoration. It uses a Sobel filter for edges, median filters for noise reduction, and pixel shuffling for efficient upsampling, converting LR images to HR while leveraging high-dimensional features.

2.2. Lightweight image SR methods

Lightweight SR research increasingly aims to balance reconstruction quality with low computational cost. Li *et al.* [18] introduced FourierSR, which reduces complexity using frequency-domain operations but struggles to capture fine local textures. Li *et al.* [19] proposed lightweight cross-receptive focused inference network (CFIN), blending CNN and Transformer features for improved contextual modeling, though its multiple modules increase architectural overhead. Gao *et al.* [20] developed FDIWN with feature distillation and weighting for efficient reuse, but its numerous distillation blocks make optimization more difficult. Xue *et al.* [21] designed multi-path feedback fusion network (MFFN) to enhance iterative refinement using fusion-attention feedback blocks, yet its multi-path structure raises memory usage and slows inference. Li *et al.* [22] proposed a lightweight adaptive weighted super-resolution network (LW-AWSRN) that integrates fusion local fusion block (LFB) with adaptive multi-scale, but it still adds extra components to manage scale redundancy. However, the use of multiple scale branches still introduces parameter redundancy, limiting its suitability for highly resource-constrained deployments. However, MFFN, and LFB-based networks improve reconstruction through iterative or multi-path refinement, but each suffers from drawbacks that limit lightweight deployment. In contrast, the proposed LPFN introduces a streamlined parallel feedback mechanism that reuses features more efficiently.

2.3. Reinforcement learning-based SR methods

In this study, Chen *et al.* [23] proposed spatial-temporal hierarchical reinforcement learning (STAR-RL) for pathology image SR. It models SR as an MDP, using a hierarchical patch-level recovery with a spatial manager to target degraded patches and a temporal manager to decide early stopping, improving reconstruction efficiency and avoiding over-processing. In the context of single image SR, Bouffard *et al.* [24] proposed a multi-agent reinforcement learning (MRL) approach for single-image SR, where agents adjust pixel intensities using local enhancement operators in a content-aware manner. This method improves resolution without the complexity of GAN-based model. Altinkaya and Barakli [25] proposed (DRL-SRFR), a deep reinforcement learning based SR of face regions. It uses visual attention and RRDBs for iterative, patch-based detail enhancement, but its use of DRL and dense residual blocks increases computational complexity compared to traditional methods.

3. PROPOSED METHODOLOGY

The LPFN model follows a clear end-to-end pipeline where each component refines the features passed from the previous stage. First, the LR image is processed by the feedback block, which extracts initial features while also receiving high-level feedback from later layers to correct early feature errors. The refined feedback block output is then sent to the dispersion-aware attention residual block (DARB), where spatial and channel attention selectively enhances important textures and edges. Next, the enhanced features move into EdgeNet, a lightweight module that sharpens boundaries and fine structures before upsampling. During training, the CRL module supervises all stages through a global feedback loss, which adjusts automatically based on the learning progress. This allows feedback block, DARB, and EdgeNet to operate in one unified feedback pathway, enabling stable training, faster convergence, and clearer reconstruction.

3.1. Lightweight parallel feedback network (LPFN)

LPFN is designed for high-quality SR on devices with limited computing power. It takes a LR image as input and improves it step-by-step using a parallel feedback mechanism that reuses and refines features, helping the network correct earlier feature errors for clearer textures and sharper edges. Figure 1 presents the overall architecture of the proposed LPFN model, which performs SR through a progressive feature-refinement pipeline. The LR image first undergoes shallow feature extraction, and the resulting features are passed into the feedback blocks. These blocks iteratively refine low-level features by using feedback signals from deeper layers, helping the model correct early feature errors. The refined features are then sent to the DARB, where channel and spatial attention emphasize important textures, edges, and structural information. The enhanced features are further processed by EdgeNet, a lightweight edge-sharpening unit designed to strengthen boundary details before upsampling. All improved features are

finally fused and upsampled to generate the HR output. During training, the CRL module supervises the entire architecture using a global feedback loss, ensuring that feedback block, DARB, and EdgeNet work smoothly together during training to produce clearer details and more accurate SR outcome. The outcome of the part on shallow feature extraction is indicated as l_s , which can be calculated as follows,

$$l_s = F_c (IR) \quad (1)$$

To extract shallow low R features, F_c consist of (3,128) and (128, 32) convolutions as in (1). In the feedback part set, n is considered as several feedbacks from 1 to N and t is considered as several iterations from 1 to T . Hence the output of the n -th term of the feedback network in t -th iteration is defined as, l_n^t . In the first iteration, the first feedback input and other outputs contain only shallow features l_s , while residual learning allows subsequent feedback outputs to include results from previous networks $l_1^1 \dots l_{n-1}^1$ as described in (2),

$$l_n^1 = \begin{cases} F_{\text{feedback}}(l_s) & n = 1 \\ F_{\text{feedback}}(l_s, l_1^1 \dots l_{n-1}^1) & n \geq 2 \end{cases} \quad (2)$$

where operations of feedback block are considered as F_{feedback} . In the reconstruction stage, the concatenated feedback block outputs are upsampled via deconvolutional layers,

In the global feedback stage, the downsampling operator generates IR' in (3), which is compared with the original IR to supervise and guide the training of the LPFN.

$$IR' = F_{\text{downsample}}(\text{Super} - \text{Resolution}), \quad (3)$$

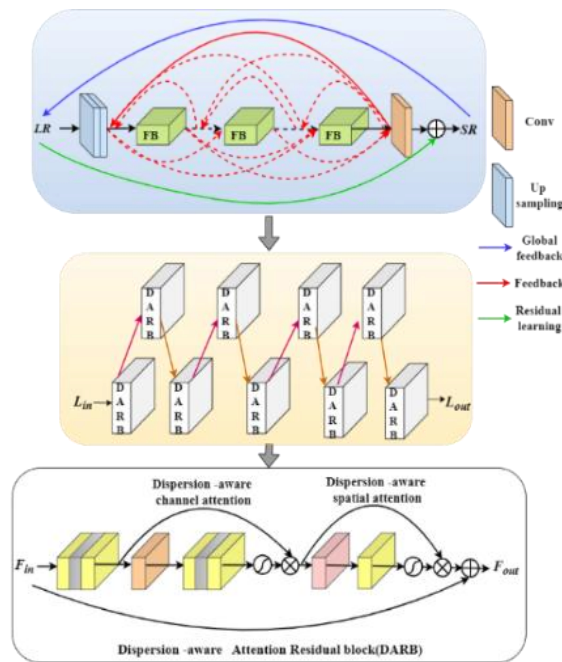


Figure 1. Proposed LPFN framework

3.2. Feedback block

After the LPFN module generates the refined feature output, it is sent back into the feedback block to begin the next refinement cycle, where low-level features are iteratively improved using high-level information, resulting in enhanced edge sharpness and texture reconstruction. Through repeated upsampling and downsampling of HR and LR feature flows defined as h_1, h_2, h_3, h_4 and $l_{in}, l_1, l_2, l_3, l_4, l_{out}$ respectively, it enhances convergence, stability, and efficient feature reuse, boosting reconstruction performance while keeping computational cost low, making it suitable for resource-constrained environments. In the feedback block, set the number of projections as $g = 4$. The feedback process is defined as follows,

$$l_g = \begin{cases} F_{\text{DARB}}(l_{\text{in}}), & g = 1 \\ F_{\text{DARB}}(F_{\text{downsample}}(h_{g-1})), & 2 \leq g \leq 4 \end{cases}, \quad (4)$$

$$h_g = F_{\text{DARB}}(F_{\text{upsample}}(l_g)) \quad 1 \leq g \leq 4, \quad (5)$$

$$l_{\text{out}} = F_{\text{DARB}}(F_{\text{downsample}}(h_4)), \quad (6)$$

In (4)-(6), F_{DARB} is considered as a basic block of DARB in LR and HR feature flows and F_{upsample} , $F_{\text{downsample}}$ denotes the processes for upsampling using deconvolution and downsampling through convolution, respectively.

3.3. Dispersion-aware attention residual block (DARB)

The features refined by the feedback block are then passed to the DARB module, to enhance the textures and edge regions. The DARB, shown in Figure 2, to improve the efficiency of the feedback block. DARB enhances feature refinement in LPFN by integrating dispersion-aware channel attention (DACA) and spatial attention (DASA) within a residual structure. It uses standard deviation (SD) alongside average and max pooling to capture rich representations, with channel attention emphasizing important channels and spatial attention refining pixel-level details. When combined with feedback blocks, DARB enables iterative enhancement of edges and textures, improving SR quality with low computational cost.

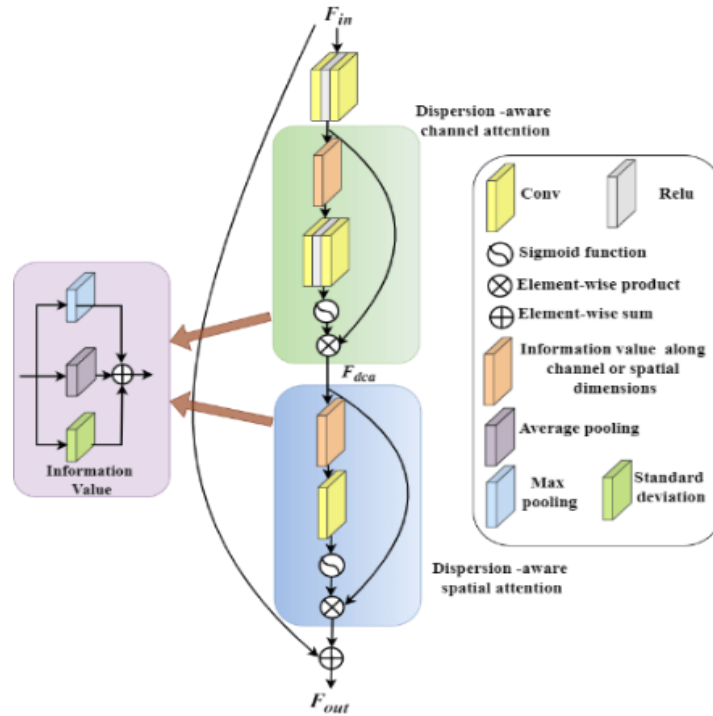


Figure 2. Dispersion-aware attention residual block (DARB)

3.3.1. Dispersion-aware channel attention (DACA)

To enrich channel representation, SD, average pooling, and max pooling descriptors are combined. Average pooling summarizes overall channel information, while max pooling highlights prominent features. SD adds a novel perspective by capturing pixel dispersion and structural boundaries. Together, they provide a richer context, enabling DACA mechanism to significantly boost the discriminative power of feature maps.

A 1D channel attention map $R^{C \times 1 \times 1}$ is generated using an MLP and sigmoid normalization, then applied via element-wise multiplication to the input feature map f_{in} . The final DACA output is defined in (7).

$$f_{\text{dca}} = f_{\text{in}} * \sigma F_{\text{mip}}(V_c). \quad (7)$$

3.3.2. Dispersion-aware spatial attention (DASA)

Similar to the DACA, DASA enhances spatial feature refinement by combining those descriptors. While average and max pooling capture overall and prominent pixel values, standard deviation highlights pixel dispersion across channels, allowing the network to focus on crucial spatial information.

A 2D spatial attention map $R^{1 \times h \times w}$ is produced using a 7×7 convolution and sigmoid activation, then applied to the input f_{dca} via element-wise multiplication [25]. The final DASA output is given in (8):

$$f_{dsa} = f_{dca} * \sigma \left(f^{7 \times 7}(V_{i,j}) \right). \quad (8)$$

Finally, f_{out} , as the output of DARB, can be obtained by using (9):

$$f_{out} = f_{dca} + f_{in}. \quad (9)$$

3.4. EdgeNet

EdgeNet, integrated into the LPFN framework, enhances edge clarity and texture sharpness in HR images with low computational cost. Inspired by richer convolutional features (RCF) but simplified from five stages to three and using depth-wise separable convolutions, it reduces parameters while preserving performance. By focusing on edge-specific features, EdgeNet prevents blurring, improves SR visual quality, and remains efficient for resource-constrained environments.

Figure 3 illustrates EdgeNet, a lightweight adaptation of the RCF edge detection model. EdgeNet reduces the original five-stage VGG16 design to three stages and replaces 3×3 convolutions with depth-wise separable convolutions ('Disconv') for efficiency. It uses deconvolution ('Deconv') and pooling layers for feature reconstruction and reduction, with kernel notation " $k \times k - m$ " indicating size and filter count. This streamlined architecture efficiently extracts and reconstructs edge features while reducing computational cost, making it effective for edge enhancement in SR.

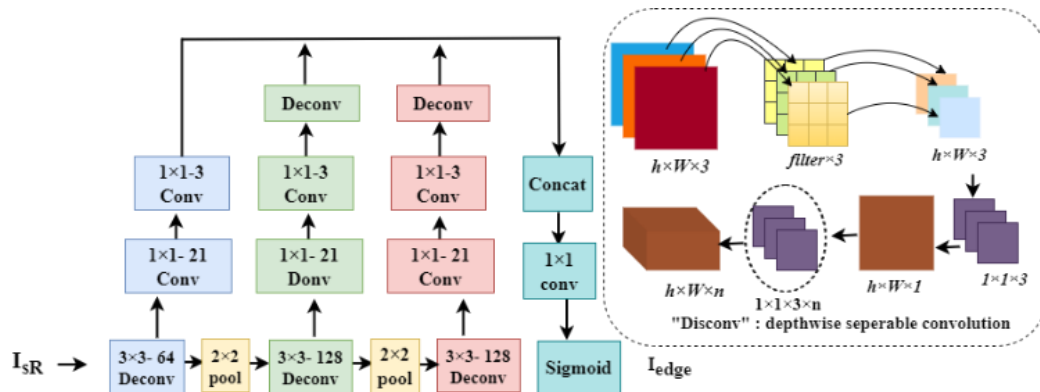


Figure 3. The framework of edge enhanced network (EdgeNet)

3.4.1. Fusion

The fusion mechanism combines the super-resolved image I_{SR} and edge-enhanced image I_{Edge} to integrate texture and edge details, enhancing clarity and sharpness. The fusion mechanism combines the super-resolved image (I_{SR}) and edge-enhanced image (I_{Edge}) through concatenation, followed by a 1×1 convolution to reduce dimensionality and generate the final output. This process integrates enhanced edge details with texture information, producing sharper and clearer reconstructions. Additionally, the mechanism leverages information from projection groups to refine LR features for subsequent iterations, enabling effective feedback generation and improved reconstruction quality.

3.5. Global feedback loss function

LPFN uses a global feedback mechanism to relay SR degradation information to LR images, improving LR-HR mapping. It adaptively balances image quality and efficiency using reinforcement learning, with two losses: the primary SR-HR regression loss and a feedback-regression loss for LR and modified LR', defined as:

$$\text{loss} = l_1(\text{SR}, \text{hR}) + \theta l_1(\text{LR}', \text{LR}), \quad (10)$$

where θ , controls the feedback loss weight in (10). Feedback is applied at the network's end, optimized via curriculum reinforcement learning, and generated with lightweight convolutional downsampling and channel-transform layers for minimal parameter overhead.

3.6. CRL with policy transfer and enhancement

The proposed CRL-based method uses reinforcement learning to adaptively update weights and optimize global feedback loss for improved SR performance. Using the l_1 loss, the LPFN adapts to varying image complexities by arranging HR targets sequentially over T iterations, following a curriculum strategy. The CRL framework progressively refines image features through a neural policy that generates optimal feature refinement sequences, maximizing a reward based on pixel-level accuracy and perceptual quality for effective SR enhancement. In (11) formulates the policy optimization objective as maximizing the overall SR reward to achieve the most effective feature enhancement strategy.

$$\max_{\theta} R(\xi^*(z(\theta))) \text{ s.t. } \xi^*(z(\theta)) = f_{\text{MPC}}(z(\theta)) \quad (11)$$

The reward R is calculated with respect to network parameters θ using the chain rule for gradient computation.

$$\frac{dR}{d\theta} = \frac{\partial r}{\partial z} \frac{\partial z}{\partial \theta}. \quad (12)$$

Here, $\frac{\partial z}{\partial \theta}$ is automatically obtained through backpropagation, while $\frac{dR}{d\theta}$ accounts for the SR-specific reward gradients associated with reconstruction quality. The on-policy RL framework directly optimizes network parameters using SR-relevant rewards, enabling stable and efficient learning. The CRL framework mitigates this by progressively shaping the learning process through curriculum-based policy transfer, which accelerates exploration, reduces instability, and enhances feature refinement across different SR difficulty levels. Three curriculum modes $\mathcal{C} = \{c_i\}$, $i \in \{1, 2, 3\}$ are introduced to systematically transfer and refine the learned policy.

- Curriculum 1: Reward shaping for transferable SR feature refinement

This stage focuses on simple LR images with mild degradation, allowing the model to learn basic feature enhancement such as edge sharpening and structure preservation. The reward function guides convergence of essential SR parameters, helping the network develop a transferable enhancement strategy for future, more complex stages.

- Curriculum 2: Learning super resolution Policy under moderate degradation conditions

In this stage, the pre-trained policy from Curriculum 1 is refined on images with moderate degradation, including noise, blur, and compression artifacts. The curriculum enhances the model's ability to generalize under more challenging SR conditions, improving both structural and textural restoration.

- Curriculum 3: Enhancing super-resolution Policy under complex degradation

The final stage targets severely degraded images with heavy blur and fine-detail loss. The reward function emphasizes pixel-level perceptual quality and penalizes visual artifacts. Starting from the policy learned in Curriculum 2, the network progressively refines its reconstruction capability, converging to an optimal SR policy that handles diverse degradations while preserving visual fidelity.

4. RESULTS AND DISCUSSION

This section outlines the implementation and evaluation details of the proposed method. Evaluation is conducted using standard metrics, including PSNR and SSIM, which collectively assess the accuracy, visual quality, and structural fidelity of the super-resolved images.

4.1. Implementation details

The proposed architecture uses PReLU activations to process 40×40 LR RGB patches paired with HR images. Implemented in PyTorch on an NVIDIA 2070Ti GPU, it is trained with a batch size of 16 using l_1 loss and the Adam optimizer for 1,000 epochs. To maintain a lightweight design, the model employs 2 shared-parameter feedbacks while retaining competitive reconstruction capability. During training, convolutional layers use 64 filters, with kernel size and stride adjusted by upsampling scale $2 \times$, $k = 6$ and $s = 2$; for scales of $3 \times$ and $4 \times$, $k = 3$ and $s = 1$. This adaptive configuration is particularly important for lightweight SR models, helping avoid the performance degradation sometimes reported by deeper CNN-based SR approaches.

4.2. Dataset description

The proposed model is evaluated on two benchmark datasets commonly used for image SR: DIV2K and Flickr2K. DIV2K consists of 1,000 HR RGB images with diverse content, split into 800 training, 100 validation, and 100 test images. LR counterparts are generated using degradations such as bicubic downsampling, motion blur, poisson noise, and pixel shifting, with training images downsampled at factors of 0.6–0.9 and rotated to create 16,000 HR-LR image pairs. This dataset setup ensures comparability with many state-of-the-art methods that also rely on DIV2K’s standardized degradation settings. Flickr2K contains 2,650 HR images sourced from Flickr, featuring varied content and quality, including landscapes, portraits, and still-life photography, providing additional diversity and challenges for SR evaluation.

Table 1 shows that our proposed LPFN achieves optimal $\times 4$ SR performance on Set5, Set14, B100, Urban100, and Manga109 when the feedback-regression loss weight (θ) is 0.1, with PSNR and SSIM values higher than for other θ settings. Increasing θ beyond 0.1 leads to lower PSNR and SSIM, indicating diminished reconstruction quality due to excessive feedback weighting. Therefore, the ideal balance for achieving high-quality image SR is found at $\theta=0.1$. This trend aligns with observations in earlier SR literature, where overly strong feedback or attention weighting has been linked to instability and overshooting during optimization. These findings suggest that LPFN benefits from moderate feedback reinforcement without adversely affecting structural fidelity.

Table 1. Comparative analysis of feedback-regression loss weights in LPFN across various datasets

Weight	Scale	Parameters	Set5 PSNR/SSIM	Set14 PSNR/SSIM	B100 PSNR/SSIM	Urban100 PSNR/SSIM	Manga109 PSNR/SSIM
$\theta = 0$	$\times 4$	649K	35.54/0.9350	31.42/0.8652	30.20/0.8142	29.10/0.8614	34.54/0.9281
$\theta = 0.01$			35.55/0.9354	31.43/0.8654	30.22/0.8148	29.12/0.8622	34.55/0.9283
$\theta = 0.1$			35.56/0.9356	31.50/0.8663	30.25/0.8153	29.15/0.8631	34.56/0.9285
$\theta = 0.2$			35.54/0.9353	31.47/0.8660	30.24/0.8150	29.14/0.8625	34.54/0.9284
$\theta = 1$			34.46/0.9342	31.36/0.8628	30.20/0.8142	29.10/0.8608	34.46/0.9275

Figure 4 illustrates the progression of reward optimization across various RL models including the proposed CRL method, when applied to curriculum-based reinforcement learning for image SR. The reward increases over iterations, with the proposed method outperforming other models in both convergence rate and final reward level, highlighting its efficiency. The improvement highlights the effectiveness of CRL in the proposed model, which progressively adjusts task difficulty to enhance learning, optimize policies efficiently, and achieve higher rewards, especially in resource-constrained image processing tasks. Compared with existing RL approaches such as MRL, DRL-SRFR, STAR-RL, and DRL, our results demonstrate that achieving higher reward values does not come at the expense of visual quality. This contrasts with several RL-based SR models, where aggressively maximizing reward signals often leads to soft, texture-diminished, or overly smooth outputs.

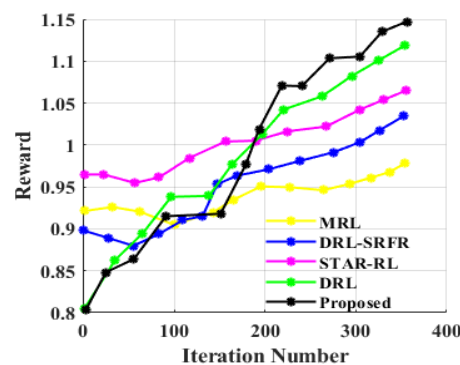


Figure 4. Comparison of CRL models in reward optimization

Figure 5 illustrates the qualitative comparison of EdgeNet with SeaNet, DASRNet, MemNet, and SR DenseNet on the DIV2K and Flickr2K datasets. The proposed EdgeNet provides visibly cleaner and sharper reconstructions, especially in challenging regions such as text boundaries, object contours, and fine

texture patterns. While SeaNet and DASRNet often generate slightly blurred edges and miss fine details, MemNet and SR DenseNet tend to lose textures and produce soft outputs. In contrast, EdgeNet keeps thin lines, character strokes, and object boundaries much sharper, with fewer artifacts and less blurring. The enhanced clarity in character strokes, boundary transitions, and high-frequency textures demonstrates the effectiveness of EdgeNet's edge extraction and refinement mechanism. This shows that EdgeNet preserves important structural details more effectively than the other methods in difficult regions, making it more reliable for high-detail SR tasks.

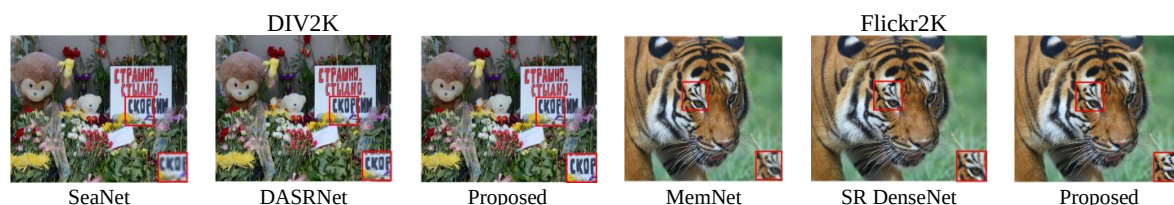


Figure 5. Visual comparisons for EdgeNet

Figure 6 presents a comparative analysis of image SR techniques on selected patches from the Flickr2K and DIV2K datasets, focusing on the effectiveness of each method in reconstructing fine image details. Techniques such as Bicubic, FISRCN, CCFN, LW-AWSRN, and EMASRN provide gradual improvements, yet often fail to recover fine details or maintain structural consistency under limited computational budgets. However, MFFN, rely on deep fusion branches but lack dynamic loss adaptation. In contrast the LPFN framework demonstrates clearer textures and sharper edges, and improve the reconstruction quality without increasing model size. These shows that high-frequency recovery benefits more from effective feature reuse and stable optimization than from deeper architectures. Overall, the findings highlight that lightweight SR models can deliver high visual quality when supported by efficient feedback mechanisms and adaptive learning strategies, offering a promising direction in SR applications.

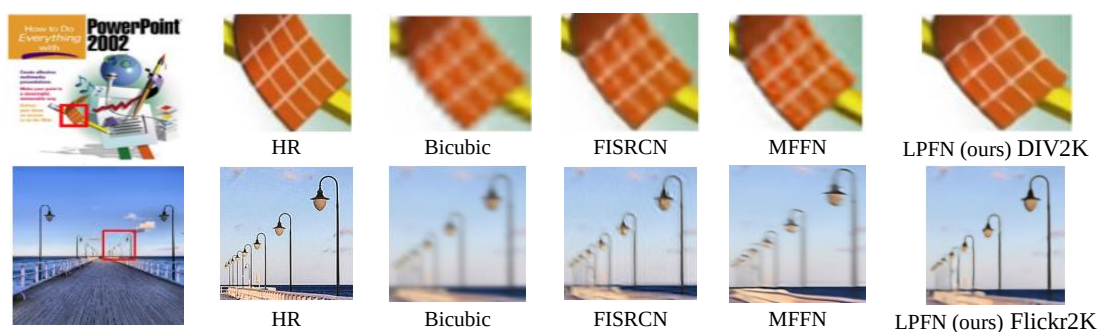


Figure 6. Visual comparisons for image SR techniques on Flickr2K and DIV2K datasets

5. DISCUSSION

Beyond visual quality improvements, LPFN has strong relevance to information and communication technology (ICT) applications. In communication systems, it can be directly integrated into image and video transmission pipelines, improving the quality of transmitted LR images and videos, reducing bandwidth requirements while preserving critical visual details. In telemedicine, LPFN can assist doctors by enhancing diagnostic images captured under low bandwidth or noisy environments, enabling more reliable remote consultations. Similarly, in network-based surveillance, the model can reconstruct clearer faces, objects, and scene details from LR camera feeds, improving recognition accuracy and situational awareness. These application-oriented integrations strengthen its practical relevance and demonstrate how LPFN can function as an adaptable enhancement module across various ICT platforms. These integration possibilities highlight the practical value of LPFN in ICT environments that demand efficient, reliable, and high-quality image enhancement under constrained computational resources.

6. CONCLUSION

This work aimed to improve lightweight SR methods by developing an efficient DL framework capable of producing high-quality reconstructions with lower computational cost than existing systems. Recent evaluations indicate that achieving high-quality SR in lightweight settings remains challenging due to loss of texture and unstable reconstruction. By integrating LPFN with feedback refinement through feedback blocks, selective feature enhancement using DARB, and edge sharpening with EdgeNet, supported by a CRL-driven global feedback loss for stable optimization. The results obtained consistently show that the proposed framework improves structural consistency and enhances edge and texture clarity even under strict computational limits. These improvements highlight that the performance gain is primarily driven by effective feature refinement and stable optimization rather than increased model size or deeper architectures. Overall, the findings confirm that lightweight SR can achieve strong visual quality when supported by efficient feedback mechanisms and adaptive learning processes in real-world SR applications. Future studies may explore extending LPFN to handle real-world degradations and video SR, with feasible approaches for integrating transformer-based modules, as well as investigating deployment on mobile/IoT devices and applications in medical and satellite imaging.

ACKNOWLEDGMENTS

All listed authors have made substantial contributions to the research, have read and approved the manuscript, and agree with its content and submission.

FUNDING INFORMATION

This research did not receive any specific grant from funding agencies in the public, commercial, or not-for-profit sectors.

AUTHOR CONTRIBUTIONS STATEMENT

This journal uses the Contributor Roles Taxonomy (CRediT) to recognize individual author contributions, reduce authorship disputes, and facilitate collaboration.

Name of Author	C	M	So	Va	Fo	I	R	D	O	E	Vi	Su	P	Fu
Srikurmam V. R.	✓			✓	✓	✓	✓		✓	✓				
Manimala														
Tadikonda Kavitha	✓	✓	✓			✓		✓	✓	✓	✓	✓	✓	

C : Conceptualization	I : Investigation	Vi : Visualization
M : Methodology	R : Resources	Su : Supervision
So : Software	D : Data Curation	P : Project administration
Va : Validation	O : Writing - Original Draft	Fu : Funding acquisition
Fo : Formal analysis	E : Writing - Review & Editing	

CONFLICT OF INTEREST STATEMENT

The authors declare that there are no conflicts of interest associated with this manuscript.

DATA AVAILABILITY

The data supporting this study are available upon reasonable request.

REFERENCES

[1] W. Li *et al.*, “Efficient face super-resolution via wavelet-based feature enhancement network,” in *Proceedings of the 32nd ACM International Conference on Multimedia*, Oct. 2024, pp. 4515–4523, doi: 10.1145/3664647.3681088.

[2] D. C. Lepcha, B. Goyal, A. Dogra, and V. Goyal, “Image super-resolution: a comprehensive review, recent trends, challenges and applications,” *Information Fusion*, vol. 91, pp. 230–260, Mar. 2023, doi: 10.1016/j.inffus.2022.10.007.

[3] W. Li, H. Guo, Y. Hou, G. Gao, and Z. Ma, “Dual-domain modulation network for lightweight image super-resolution,” *IEEE Transactions on Multimedia*, pp. 1–11, 2026, doi: 10.1109/tmm.2026.3660149.

[4] H. Yan, Z. Wang, Z. Xu, Z. Wang, Z. Wu, and R. Lyu, “Research on image super-resolution reconstruction mechanism based on convolutional neural network,” in *International Conference on Artificial Intelligence, Automation and High Performance Computing*, 2024, pp. 142–146, doi: 10.1145/3690931.3690956.

[5] J. Liang, J. Cao, G. Sun, K. Zhang, L. Van Gool, and R. Timofte, “SwinIR: image restoration using swin transformer,” in *2021 IEEE/CVF International Conference on Computer Vision Workshops (ICCVW)*, Oct. 2021, pp. 1833–1844, doi: 10.1109/iccvw54120.2021.00210.

- [6] C. Tian, Y. Yuan, S. Zhang, C.-W. Lin, W. Zuo, and D. Zhang, "Image super-resolution with an enhanced group convolutional neural network," *Neural Networks*, vol. 153, pp. 373–385, 2022, doi: 10.1016/j.neunet.2022.06.009.
- [7] N. P. M. K. N. N. S. M. S. R. and S. K. V., "Perceptual image super resolution using deep learning and super resolution convolution neural networks (SRCNN)," in *Intelligent Systems and Computer Technology*, IOS Press, 2020.
- [8] S. M. A. Bashir, Y. Wang, M. Khan, and Y. Niu, "A comprehensive review of deep learning-based single image super-resolution," *PeerJ Computer Science*, vol. 7, p. e621, 2021, doi: 10.7717/peerj-cs.621.
- [9] X. Deng, Y. Zhang, M. Xu, S. Gu, and Y. Duan, "Deep coupled feedback network for joint exposure fusion and image super-resolution," *IEEE Transactions on Image Processing*, vol. 30, pp. 3098–3112, 2021, doi: 10.1109/tip.2021.3058764.
- [10] H. Guo, J. Li, T. Dai, Z. Ouyang, X. Ren, and S.-T. Xia, "Mambair: a simple baseline for image restoration with state-space model," in *Computer Vision – ECCV 2024*, Springer Nature Switzerland, 2024, pp. 222–241.
- [11] B. Wang, B. Yan, C. Liu, R. Hwangbo, G. Jeon, and X. Yang, "Lightweight bidirectional feedback network for image super-resolution," *Computers and Electrical Engineering*, vol. 102, p. 108254, 2022, doi: 10.1016/j.compeleceng.2022.108254.
- [12] Z. Lin *et al.*, "Revisiting rcan: improved training for image super-resolution," arXiv preprint arXiv:2201.11279, 2022.
- [13] W. Li *et al.*, "Efficient image super-resolution with feature interaction weighted hybrid network," *IEEE Transactions on Multimedia*, vol. 27, pp. 2256–2267, 2025, doi: 10.1109/tmm.2024.3521753.
- [14] Y. Xiong *et al.*, "Improved SRGAN for remote sensing image super-resolution across locations and sensors," *Remote Sensing*, vol. 12, no. 8, p. 1263, Apr. 2020, doi: 10.3390/rs12081263.
- [15] F. Fang, J. Li, and T. Zeng, "Soft-edge assisted network for single image super-resolution," *IEEE Transactions on Image Processing*, vol. 29, pp. 4656–4668, 2020, doi: 10.1109/tip.2020.2973769.
- [16] M. Zhao, X. Liu, X. Yao, and K. He, "Better visual image super-resolution with laplacian pyramid of generative adversarial networks," *Computers, Materials & Continua*, vol. 64, no. 3, pp. 1601–1614, 2020, doi: 10.32604/cmc.2020.09754.
- [17] L. Kong, F. Wang, F. Yang, L. Leng, and H. Zhang, "FISRCN: a single small-sized image super-resolution convolutional neural network by using edge detection," *Multimedia Tools and Applications*, vol. 83, no. 7, pp. 19609–19627, 2023, doi: 10.1007/s11042-023-15380-3.
- [18] W. Li, H. Guo, Y. Hou, and Z. Ma, "FourierSR: a fourier token-based plugin for efficient image super-resolution," *IEEE Transactions on Image Processing*, vol. 35, pp. 732–742, 2026, doi: 10.1109/tip.2025.3648872.
- [19] W. Li *et al.*, "Cross-receptive focused inference network for lightweight image super-resolution," *IEEE Transactions on Multimedia*, vol. 26, pp. 864–877, 2024, doi: 10.1109/tmm.2023.3272474.
- [20] G. Gao, W. Li, J. Li, F. Wu, H. Lu, and Y. Yu, "Feature distillation interaction weighting network for lightweight image super-resolution," *Proceedings of the AAAI Conference on Artificial Intelligence*, vol. 36, no. 1, pp. 661–669, 2022, doi: 10.1609/aaai.v36i1.19946.
- [21] L. Xue, J. Shen, R. Wang, and J. Yang, "MFFN: multi-path feedback fusion network for lightweight image super resolution," *IET Image Processing*, vol. 17, no. 14, pp. 4190–4201, 2023, doi: 10.1049/ipr2.12927.
- [22] Z. Li, C. Wang, J. Wang, S. Ying, and J. Shi, "Lightweight adaptive weighted network for single image super-resolution," *Computer Vision and Image Understanding*, vol. 211, p. 103254, Oct. 2021, doi: 10.1016/j.cviu.2021.103254.
- [23] W. Chen, J. Liu, T. W. S. Chow, and Y. Yuan, "STAR-RL: spatial-temporal hierarchical reinforcement learning for interpretable pathology image super-resolution," *IEEE Transactions on Medical Imaging*, vol. 43, no. 12, pp. 4368–4379, Dec. 2024, doi: 10.1109/tmi.2024.3419809.
- [24] A. Bouffard, M. Pop, and M. Ebrahimi, "Multi-step reinforcement learning for medical image super-resolution," in *Medical Imaging 2023: Image Processing*, Apr. 2023, p. 73, doi: 10.1117/12.2653655.
- [25] E. Altinkaya and B. Barakli, "Enhancing face image quality: strategic patch selection with deep reinforcement learning and super-resolution boost via RRDB," *IEEE Access*, vol. 12, pp. 120142–120164, 2024, doi: 10.1109/access.2024.3450571.

BIOGRAPHIES OF AUTHORS



Dr. S V R Manimala    has received B.E. from the Department of Electronics and Communication, C.R.R. College of Engineering, Andhra University in 1994 and M.Tech. from JNTUH, in 2003. She completed her Ph.D. from JNTUA, Ananthapuramu in the year 2021. She has been working at MVSR Engineering College since 2000, and is currently serving as associate professor in the Department of Electronics and Communication Engineering. Her research interests include bio-medical signal processing, neural networks and machine learning, data communications. She can be contacted at email: srikurmam.manimala@gmail.com.



Dr. T Kavitha    has received B.E. from the Department of Electronics and Communication, MVSR College of Engineering, Osmania University in 1994 and M.Tech. from JNTUH, in 2023. She was awarded Ph.D. from OU, Hyderabad in the year 2023. She has been working at MVSR Engineering College since 2000, and is currently serving as associate professor in the Department of Electronics and Communication Engineering. The research areas of interest include image and video processing, communications, data compression and AI and ML. She can be contacted at email: tadikondakavitha@gmail.com.