

High-Efficiency Distillation Network for Lightweight Single-Image Super-Resolution

Xicheng Biao

School of Computer Science and Engineering

Xi'an Technological University

Xi'an, China

E-mail: biao20001009@163.com

Li Zhao

School of Computer Science and Engineering

Xi'an Technological University

Xi'an, China

E-mail: 332099732@qq.com

Abstract—Lightweight single-image super-resolution is important for edge-side imaging systems where reconstruction accuracy, model size, memory movement, and inference latency must be considered together. Although many compact SR networks reduce floating-point operations through channel distillation or depth wise-style operators, fewer FLOPs do not necessarily lead to faster execution on real hardware because intermediate feature access may still be expensive. In addition, when the distillation branch is overly compressed, shallow features may be forwarded with insufficient nonlinear transformation, which weakens the recovery of complex textures. To address these problems, this paper proposes PCF-IMDN, a compact feature-distillation network based on partial-convolution feature extraction. In the proposed block, spatial convolution is applied only to a selected subset of channels, and a following 1×1 pointwise convolution is used to fuse cross-channel information. This design reduces redundant spatial filtering while preserving feature interaction. Moreover, a lightweight 1×1 transformation is inserted into the retained distillation branch to enhance feature reuse with a small parameter increase. Batch normalization is removed to avoid disturbing low-level image statistics and to simplify inference. Experiments on DIV2K and benchmark datasets including Set5, Set14, and Urban100 show that PCF-IMDN reduces the parameter count from 715K to 430K and the computational cost from 158G to 92G compared with IMDN. The model achieves 4.75 ms latency on Set5 and 25.6 FPS for 1080P input, demonstrating its potential for real-time edge-oriented image enhancement.

Keywords- *Single-Image Super-Resolution; Lightweight Network; Partial Convolution; Feature Distillation; Edge Computing*

I. INTRODUCTION

Single-image super-resolution aims to reconstruct a high-resolution image from one low-resolution input. In practical scenarios such as mobile photography, surveillance video, remote sensing, and medical imaging, an SR model is expected not only to improve visual details but also to run within limited computational and memory budgets. Early deep-learning-based methods, including SRCNN [11], VDSR [3], EDSR [12], and RCAN [4], greatly improved reconstruction accuracy by increasing network depth, residual learning capacity, or channel attention mechanisms. However, these accuracy-oriented models usually require large numbers of parameters and heavy convolutional computation, which restricts their deployment on mobile and embedded devices.

To improve deployment efficiency, many lightweight SR models have been proposed. Representative methods such as IDN [6], IMDN, RFDN [3] RLFN, and BSRN reduce complexity through information distillation, residual reuse, local feature aggregation, or efficient residual structures. These studies show that a compact SR backbone should preserve useful hierarchical features while avoiding unnecessary computation. Nevertheless, two problems remain. First, some efficient operators reduce theoretical FLOPs but still introduce non-negligible memory access during inference. Depth wise separable convolution is a typical example: after spatial and

channel operations are separated, feature maps may need to be read and written more frequently, so the measured speedup can be smaller than expected. Second, in some ultra-lightweight distillation blocks, the retained branch is passed forward with little additional transformation. This helps reduce computation, but it may limit the model's ability to describe complex high-frequency textures.

Based on these observations, this work revisits the IMDN framework from the perspective of edge deployment. Instead of only pursuing fewer parameters, the proposed PCF-IMDN focuses on the relationship among spatial filtering, channel fusion, memory access, and feature reuse. Partial convolution is introduced into the feature extraction unit so that spatial filtering is applied only to selected channels. A 1×1 pointwise convolution then mixes information among channels. In this way, the model avoids full-channel spatial processing while retaining channel interaction. In addition, the retained distillation branch is enhanced with a lightweight 1×1 transformation, allowing distilled features to be refined before final aggregation [20]. Batch normalization is removed because SR is a low-level vision task in which preserving image statistics is important.

The main contributions of this paper are summarized as follows.

First, a partial-convolution-based feature extraction unit is designed for lightweight SR. Compared with full convolution, the proposed unit reduces redundant spatial computation; compared with depth wise-style operators, it pays more attention to feature access and hardware-friendly execution.

Second, an enhanced information distillation branch is introduced. The retained branch is no longer directly concatenated after channel splitting; instead, it is processed by a low-cost 1×1 pointwise convolution, which improves nonlinear feature transformation with limited extra parameters.

Third, a denormalized reconstruction design is adopted. By removing batch normalization layers, the network avoids possible shifts in image feature statistics and reduces inference overhead.

Finally, experiments show that PCF-IMDN achieves a better balance between reconstruction quality and deployment efficiency. Compared with IMDN, the proposed model reduces parameters from 715K to 430K and FLOPs from 158G to 92G, while maintaining competitive PSNR [18] and SSIM [21] performance.

Key contributions include:

Operator reconstruction: PConv is used to replace full-channel spatial filtering in the feature extraction block. By computing spatial convolution on only part of the channels and passing the remaining channels through a lightweight path, the block reduces redundant operations and memory traffic.

Enhanced distillation: a 1×1 pointwise convolution is added to the retained coarse feature branch, increasing transformation depth with little extra cost and alleviating the representation loss caused by compact pruning.

Denormalized reconstruction: batch-normalization layers are removed so that pixel statistics are not distorted, which benefits edge detail recovery and inference efficiency in SR tasks.

II. RELATED WORK

A. Lightweight Network Architecture

The development of lightweight SR networks is closely related to where feature extraction is performed and how intermediate features are reused. SRCNN performs feature mapping after the input image has already been enlarged, which increases the computational burden in high-resolution space. FSRCNN [30] improves this pipeline by moving most feature extraction into low-resolution space and placing the up sampling operation near the end of the network. Later high-performance models such as VDSR, EDSR, and RCAN further improve reconstruction

accuracy through deeper residual learning and attention mechanisms, but their computational cost remains high fore-edge-side deployment.

Information distillation has become an effective strategy for lightweight SR. IDN separates features into informative and remaining parts, enabling the network to reuse useful channels while reducing redundant computation. IMDN extends this idea by performing multi-stage feature distillation and aggregation, achieving a strong trade-off between parameter size and reconstruction quality. RFDN, RLFN, and BSRN further improve lightweight design by introducing residual feature distillation, local feature reuse, and blueprint-style separable residual structures. These methods indicate that compact SR models should not simply remove layers; instead, they should carefully control how features are separated, transformed, and fused.

In this paper, IMDN is selected as the baseline because its modular distillation structure is clear and suitable for operator-level optimization. The proposed PCF-IMDN keeps the general information-distillation idea but changes the internal feature extraction unit and the retained distillation branch. This allows the network to reduce computation while preserving the transformation ability needed for high-frequency details.

B. Efficient Convolution Operators

The runtime of a convolutional SR network is jointly influenced by FLOPs [12] and memory access cost (MAC). For an input feature map X in $\mathbb{R}^{H \times W \times C}$, H and W are the spatial dimensions, C is the number of channels, and k is the kernel size. A standard C-to-C convolution has a cost on the order of $H \times W \times k^2 \times C^2$, which becomes expensive when the channel width increases.

DWConv is often adopted to lower arithmetic complexity by separating spatial filtering and channel mixing. Its nominal FLOPs can be reduced to about $H \times W \times (k^2 \times C + C^2)$. Nevertheless, the separated stages require additional intermediate feature storage and retrieval. When

channels are widened to recover accuracy, the MAC term in Eq. (1) can dominate actual latency.

$$MAC_{DW} = h \times w \times 2c' + k^2 \times c \quad (1)$$

Therefore, a design with fewer FLOPs may still run slowly if it repeatedly moves feature maps through memory.

Here, MAC characterizes the amount of feature data read from and written to memory; H , W , C , and B denote feature-map size, channel width, and bytes per activation. This metric explains why operator-level latency cannot be evaluated by FLOPs alone [28].

PConv provides a more hardware-aware alternative. It applies $k \times k$ spatial convolution only to a channel subset $C_p = C / n_{div}$ and leaves the other $C - C_p$ channels on an identity path. Local spatial responses are thus extracted without repeatedly processing every channel [33]. Under this division, the spatial branch costs about $H \times W \times k^2 \times C_p^2$, and the simplified access pattern is summarized in Eq. (2):

$$\begin{aligned} MAC_{PConv} &= h \times w \times 2c_p + k^2 \times c_p^2 \\ &\approx \frac{1}{4} \times MAC_{standard} \end{aligned} \quad (2)$$

The proposed network embeds this operator into SR feature extraction to reduce the gap between theoretical complexity and measured speed. By decreasing unnecessary memory movement, the model improves edge throughput while keeping reconstruction quality stable.

The complete inference path is: LR input $\rightarrow$ shallow feature extraction $\rightarrow$ stacked PCF-IMDB modules with PConv/PWConv and enhanced distillation $\rightarrow$ global feature fusion $\rightarrow$ sub-pixel reconstruction $\rightarrow$ SR output. The training objective in Eq. (3) constrains the final SR output using the HR image.

C. Feature Distillation and Pointwise Convolution

Feature distillation aims to pass useful information forward while progressively refining representation granularity. In IDN, IMDN, and related compact networks, PWConv is commonly used to exchange channel information with a small parameter budget, making it suitable for lightweight SR backbones.

However, if the retained branch in a distillation block is simply forwarded, the branch contributes little nonlinear transformation. This reduces computation but may weaken the modeling of shallow features. Since spatial convolution alone cannot sufficiently mix inter-channel information, combining PConv with PWConv can approach the representation ability of standard convolution while preserving efficiency.

III. METHODOLOGY

A. Overall Network Architecture

Fig. 1 shows the PCF-IMDN architecture. The network follows the IMDN pipeline and contains three parts: a shallow feature extractor, a nonlinear mapping stage composed of stacked PCF-IMDB modules, and a reconstruction head.

Shallow Feature Extraction: the LR input is first mapped into shallow feature space by an initial convolution.

Nonlinear Mapping: six PCF-IMDB modules are cascaded to transform and refine deep features.

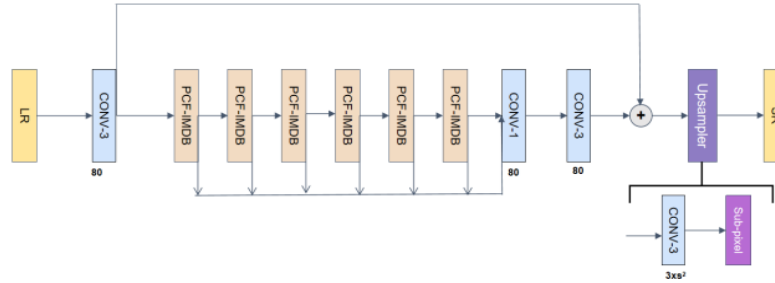


Figure 1. Architecture of the proposed PCF-IMDN network.

Reconstruction module: outputs from the mapping stage are fused, compressed by convolution, and up sampled through sub-pixel reconstruction to obtain the HR image.

B. Enhanced Information Multi-Distillation Module

The conventional IMDB is compact, but its internal convolution and retained branches can still be optimized. PCF-IMDB is therefore designed with two targeted modifications.

1) Partial-convolution-based feature extraction

In the feature extraction unit shown in Fig.2, PConv is used to reduce the cost of full-channel spatial convolution while avoiding the memory inefficiency of DWConv. Given $X \in \mathbb{R}^{(H \times W \times C)}$, channels are divided by n_{div} ; only $C_p = C / n_{div}$ channels undergo $k \times k$ spatial filtering, whereas the rest are forwarded by identity mapping. A following 1×1 PWConv fuses all channels. With $n_{div} = 4$, the module keeps enough channels for spatial detail extraction while maintaining the computational advantage of PConv.

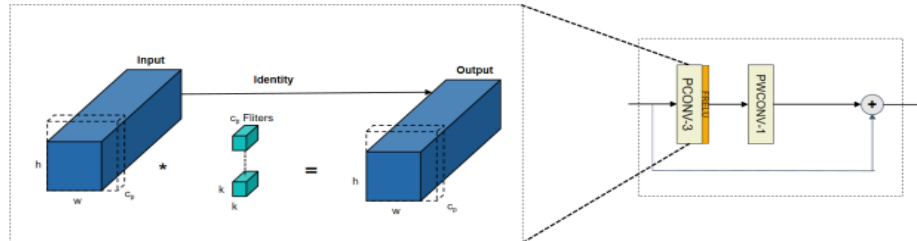


Figure 2. PConv and PWConv feature extraction unit inspired

The PConv plus PWConv cascade forms a T-shaped receptive field. Compared with full convolution, it avoids processing all channels spatially; compared with DWConv, it reduces intermediate access and improves hardware utilization.

2) Enhanced feature distillation branch

In the original IMDN block, the retained branch after channel splitting is commonly passed directly to later concatenation. For an ultra-lightweight model, such a direct path may not transform features sufficiently.

To strengthen this branch, an extra 1x1 PWConv is inserted. As illustrated in Fig. 3, the retained features are no longer concatenated immediately; instead, they pass through a lightweight transformation before fusion. This adds few parameters but provides a deeper nonlinear path for texture representation.

Within PCF-IMDB, features are divided into a distilled branch and a continuing branch. The distilled branch carries compact information, while the continuing branch keeps being transformed; concatenation at the end integrates the multi-level responses before channel fusion.

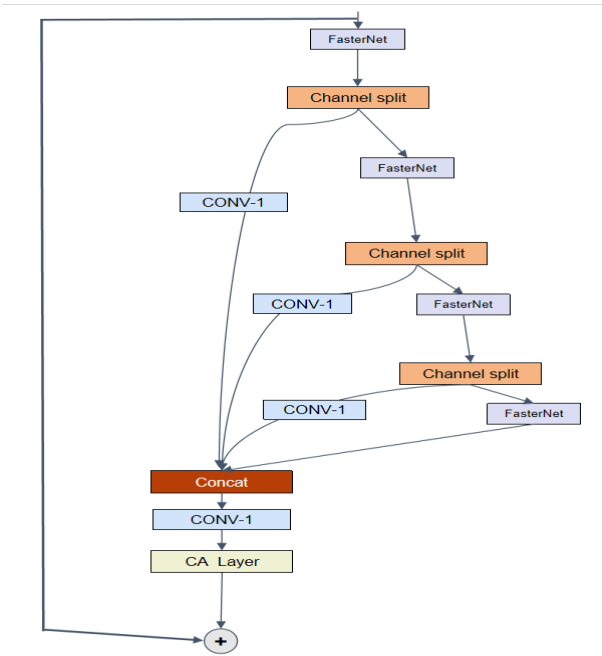


Figure 3. Structure of the proposed PCF-IMDB module.

C. Denormalization Design

BN is effective in many high-level recognition networks, but in SR it may shift image statistics and restrict the range of reconstructed features. Although BN can sometimes be fused during inference, it still adds computation and can weaken high-frequency detail. Therefore, BN is removed in PCF-IMDN to maintain image fidelity and reduce normalization overhead.

D. Loss Function

To train the model stably and compare fairly with lightweight baselines, an L1 objective is used. For N pairs $\{(I_{i\wedge LR}, I_{i\wedge HR})\}$, the network prediction is $I_{i\wedge SR} = F_{\theta}(I_{i\wedge LR})$, and the optimization target is given by Eq. (3):

$$L(\theta) = \frac{1}{N} \sum_{i=1}^N \|H_{PCF-IMDN}(I_{LR}^i) - I_{HR}^i - I_{HR}^i\|_1 \quad (3)$$

Where θ is the set of learnable parameters, N is the number of training samples, $I_{i\wedge SR}$ denotes the reconstructed image, and $I_{i\wedge HR}$ denotes the ground-truth HR image.

IV. EXPERIMENTS

A. Datasets and Evaluation Metrics

Training uses the DIV2K dataset, which provides 800 high-quality RGB images and is widely used for image restoration. Evaluation is conducted on Set5, Set14, and Urban100.

Restoration quality is measured by PSNR and SSIM. Following standard SR evaluation, metrics are computed on the Y channel in YCbCr space. Efficiency is assessed with parameters, FLOPs, latency, and frame rate.

B. Implementation Details

HR images from DIV2K are down sampled with bicubic interpolation to form LR-HR training pairs. Random crops are used as inputs, and horizontal flipping and 90-degree rotation are applied for augmentation.

All experiments are conducted in PyTorch on a single NVIDIA RTX 4060 GPU.

The ADAM optimizer is used with $\text{beta}_1 = 0.9$ and $\text{beta}_2 = 0.999$. The initial learning rate is 2×10^{-4} and follows a cosine annealing schedule. IMDN and PCF-IMDN are trained and evaluated under the same settings.

C. Model Analysis

1) Ablation Studies

To examine the contribution of each component, ablation experiments compare the original IMDN with variants that progressively add PConv, enhanced distillation, and BN removal. Table 1 reports the result.

TABLE I. PERFORMANCE AND COMPLEXITY COMPARISON BETWEEN BASELINE AND IMPROVED MODELS

Model	PConv	1x1 Enhanced Distillation	BN Layer Removal	Parameters (Params)	Computational Cost (FLOPs)	PSNR/SSIM
Baseline (IMDN)	×	×	×	715K	158G	32.21/0.8948
Model A	✓	×	×	405K	86G	31.98/0.8905
Model B	✓	✓	×	430K	92G	32.09/0.8925
Model C	✓	✓	✓	430K	92G	32.13/0.8936

Table 1 shows that PCF-IMDN becomes substantially lighter after introducing PConv, the enhanced distillation branch, and denormalization.

The parameter count decreases from 715K to 430K, reducing the model size by about 40%.

FLOPs are reduced from 158G to 92G, also corresponding to roughly a 40% reduction.

Although PSNR decreases slightly, the loss is small relative to the resource savings. The results indicate that the proposed architecture removes redundant computation while keeping reconstruction accuracy at a competitive level.

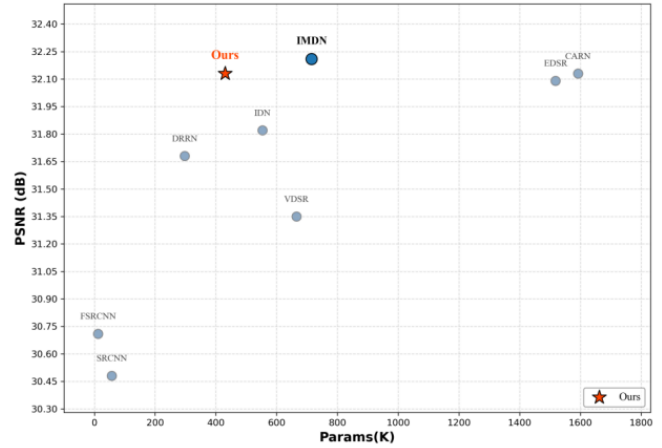


Figure 4. Trade-off between performance and parameter count on the Set5 x4 dataset.

2) Training Dynamics and Convergence Analysis

Training curves are used to verify convergence behavior. As shown in Fig. 5, the L1 loss decreases steadily, while PSNR and SSIM rise and then stabilize on the DIV2K training set.

3) Comparison with State-of-the-Art Models

PCF-IMDN is evaluated against Bicubic interpolation and lightweight SR models including SRCNN, FSRCNN, VDSR, RFDN, and IMDN. Quantitative x4 results are listed in Table 2.

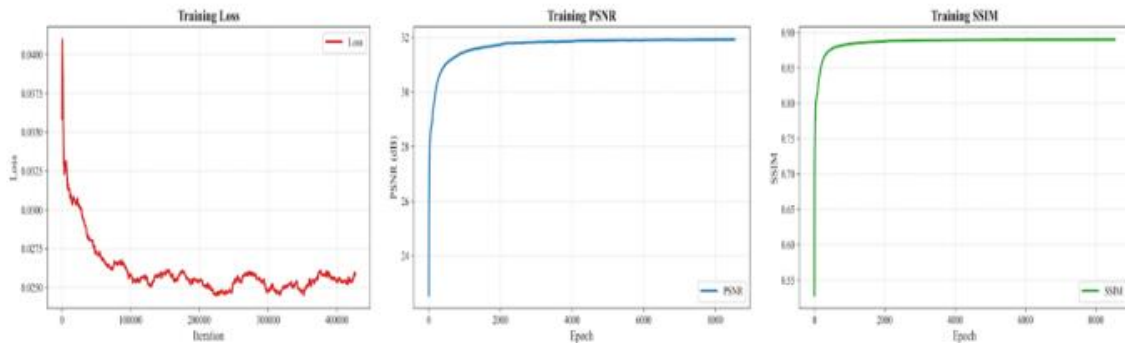


Figure 5. Convergence curves of the PCF-IMDN training process.

D. Runtime Analysis

Runtime is tested on an NVIDIA GeForce RTX 4060 Laptop GPU under multiple input resolutions, as summarized in Table 3. The scenarios range from cropped benchmark images to 720P and 1080P video streams, all with x4 super-resolution.

The results show that PConv-based memory-access optimization and BN removal improve speed in every scenario. On Set5, PCF-IMDN reduces latency to 4.75 ms. For 1080P reconstruction, it reaches 25.61 FPS, exceeding the 24 FPS real-time threshold, while IMDN obtains 21.17 FPS. This confirms the practical value of PCF-IMDN for real-time edge.

TABLE II. AVERAGE PSNR/SSIM COMPARISON OF DIFFERENT ALGORITHMS ON BENCHMARK DATASETS

Algorithm	Parameters	Set5	Set14	Urban100
		PSNR/SSIM	PSNR/SSIM	PSNR/SSIM
Bicubic	-	28.42/0.8104	26/0.7027	23.14/0.6577
SRCNN	8K	30.48/0.8628	27.50/0.7513	24.52/0.7221
FSRCNN	13K	30.72 / 0.8660	27.61 / 0.7550	24.62/0.7280
VDSR	666K	31.35 / 0.8838	28.01 / 0.7674	25.18/0.7524
RFDN	550K	32.24/0.8952	28.61/0.7819	26.11/0.7858
IMDN	715K	32.21 / 0.8948	28.58 / 0.7811	26.04/0.7838
PCF-IMDN	430K	32.13/0.8936	28.49/0.7793	25.90/0.7786

TABLE III. INFERENCE LATENCY COMPARISON UNDER DIFFERENT TEST SCENARIOS.

Test Scenario	Model	Inference Time (Latency)	Frame Rate (FPS)
Set5 (256x256)	IMDN	8.28ms	120.74
	PCF-IMDN	4.75ms	210.5
720P (1280x720)	IMDN	19.43ms	51.5
	PCF-IMDN	13.32ms	75.1
1080P (1920x1080)	IMDN	47.23ms	21.17
	PCF-IMDN	39.05ms	25.6

E. Comprehensive Performance Comparison

Additional comparisons are made with representative lightweight methods, including SRCNN, FSRCNN, and VDSR.

Fig. 6 presents visual results on architectural textures from Urban100.

With PConv and denormalized reconstruction, PCF-IMDN recovers line structures and boundaries more sharply than Bicubic interpolation. It also preserves geometric details better than several comparable methods and reduces over smoothing.

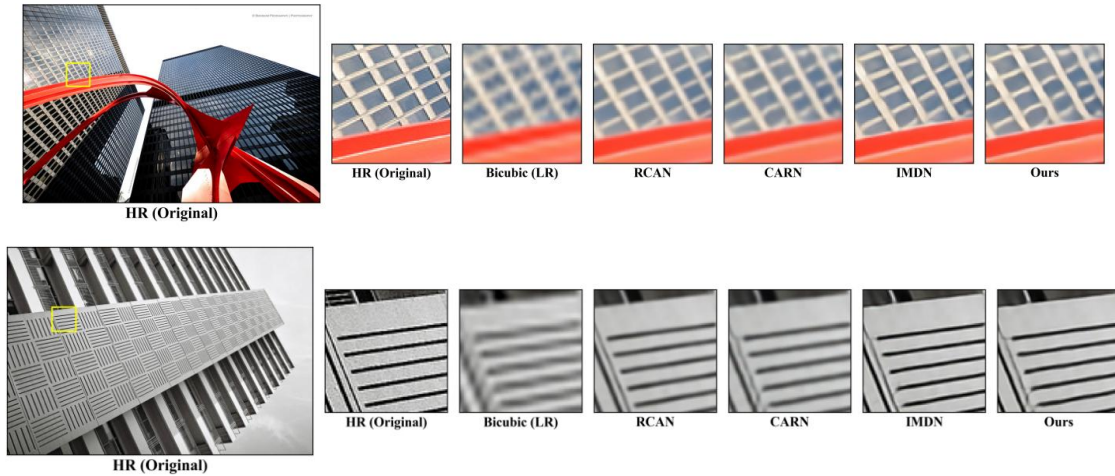


Figure 6. Visual comparison of PCF-IMDN and other SR methods on Urban100.

V. CONCLUSIONS

This work studies lightweight SR under edge-device constraints and proposes PCF-IMDN to address both computation bottlenecks and memory-access inefficiency. The design focuses on operator redesign and feature-branch enhancement.

First, the feature extraction unit is rebuilt with a PConv-PWConv cascade. Spatial convolution is performed on selected channels, while BN is removed from the improved model. This reduces redundant operations and memory access, improving hardware utilization.

Second, a 1x1 PWConv transformation is added to the distillation branch to strengthen feature representation with low overhead. The added path improves nonlinear modeling for complex textures and helps recover high-frequency information.

Experiments show that PCF-IMDN keeps competitive reconstruction quality while cutting IMDN parameters and FLOPs by about 40%. It also achieves lower latency and higher frame rate across test scenarios with acceptable PSNR and SSIM.

These findings indicate that reducing memory movement and enhancing local feature transformation are effective ways to build real-time edge-oriented SR models.

REFERENCES

- [1] C. Dong, C. C. Loy, K. He, and X. Tang, "Image super-resolution using deep convolutional networks," *IEEE Transactions on Pattern Analysis and Machine Intelligence*, vol. 38, no. 2, pp. 295–307, 2016.
- [2] J. Kim, J. K. Lee, and K. M. Lee, "Accurate image super-resolution using very deep convolutional networks," in *Proc. IEEE Conf. Comput. Vis. Pattern Recognit. (CVPR)*, 2016, pp. 1646–1654.
- [3] B. Lim, S. Son, H. Kim, S. Nah, and K. M. Lee, "Enhanced deep residual networks for single image super-resolution," in *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition Workshops (CVPRW)*, 2017, pp. 136–144.
- [4] Y. Zhang, K. Li, K. Li, L. Wang, B. Zhong, and Y. Fu, "Image super-resolution using very deep residual channel attention networks," in *Proceedings of the European Conference on Computer Vision (ECCV)*, 2018, pp. 286–301.
- [5] J. Liang, J. Cao, G. Sun, K. Zhang, L. Van Gool, and R. Timofte, "SwinIR: Image restoration using swin transformer," in *Proceedings of the IEEE/CVF International Conference on Computer Vision (ICCV)*, 2021, pp. 1833–1844.
- [6] X. Chen, J. Wang, C. Z. Zhou, and C. Dong, "Activating more pixels in image super-resolution transformer," in *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, 2023, pp. 22367–22377.
- [7] 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 *Proceedings of the European Conference on Computer Vision (ECCV)*, 2024.
- [8] H. Guo, Y. Guo, Y. Zha, Y. Zhang, W. Li, T. Dai, S.-T. Xia, and Y. Li, "MambaIRv2: Attentive state space restoration," *arXiv preprint arXiv:2411.15269*, 2024.
- [9] C. Saharia, J. Ho, W. Chan, T. Salimans, D. J. Fleet, and M. Norouzi, "Image super-resolution via iterative refinement," *IEEE Transactions on Pattern Analysis and Machine Intelligence*, vol. 45, no. 4, pp. 4713–4726, 2023.
- [10] Y. Wang, W. Yang, X. Chen, Y. Wang, L. Guo, L.-P. Chau, Z. Liu, Y. Qiao, A. C. Kot, and B. Wen, "SinSR: Diffusion-based image super-resolution in a single step," in *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, 2024.
- [11] B. Ren et al., "The tenth NTIRE 2025 efficient super-resolution challenge report," in *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition Workshops (CVPRW)*, 2025.
- [12] Z. Chen et al., "NTIRE 2025 challenge on image super-resolution (x4): Methods and results," in *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition Workshops (CVPRW)*, 2025.
- [13] B. Longarela, M. V. Conde, A. Garcia, and R. Timofte, "Efficient perceptual image super resolution: AIM 2025 study and benchmark," in *Proceedings of the IEEE/CVF International Conference on Computer Vision Workshops (ICCVW)*, 2025.
- [14] Z. Yue, K. Liao, and C. C. Loy, "Arbitrary-steps image super-resolution via diffusion inversion," in *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, 2025.
- [15] Y. A. Li, Y. Fan, X. Xiang, D. Demandolx, R. Ranjan, R. Timofte, and L. Van Gool, "Efficient and explicit modeling of image hierarchies for image restoration," in *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, 2023, pp. 18278–18289.
- [16] Z. Chen, Y. Zhang, J. Gu, L. Kong, X. Yang, and F. Yu, "Dual aggregation transformer for image super-resolution," in *Proceedings of the IEEE/CVF International Conference on Computer Vision (ICCV)*, 2023, pp. 12312–12321.
- [17] X. Wang, K. Yu, S. Wu, J. Gu, Y. Liu, C. Dong, Y. Qiao, and C. C. Loy, "ESRGAN: Enhanced super-resolution generative adversarial networks," in *Proceedings of the European Conference on Computer Vision Workshops (ECCVW)*, 2018.
- [18] X. Wang, L. Xie, C. Dong, and Y. Shan, "Real-ESRGAN: Training real-world blind super-resolution with pure synthetic data," in

- Proceedings of the IEEE/CVF International Conference on Computer Vision Workshops (ICCVW), 2021.
- [19] Z. Ji, H. Xiong, and D. Tao, "Deform-Mamba network for MRI super-resolution," in Proceedings of the International Conference on Medical Image Computing and Computer-Assisted Intervention (MICCAI), 2024.
 - [20] H. Chen, J. Gu, and Z. Liu, "OmniRWKFSR: Mamba-based super-resolution for remote sensing," arXiv preprint arXiv:2502.00404, 2025.
 - [21] C. Wang and W. Sun, "Semantic guided large scale factor remote sensing image super-resolution with generative diffusion prior," arXiv preprint arXiv:2405.07044, 2024.
 - [22] Z. Liu, S. Jiang, S. Feng, Q. Song, and J. Zhang, "Comprehensive review of deep learning approaches for single-image super-resolution," *Sensors*, vol. 25, no. 18, p. 5768, 2025.
 - [23] Y. Wen, C. Zhang, C. Xie, and X. Fu, "Achieving lightweight super-resolution for real-time computer graphics," in Proceedings of the AAAI Conference on Artificial Intelligence, 2025.
 - [24] C. Dong, C. C. Loy, K. He, and X. Tang, "Image super-resolution using deep convolutional networks," *IEEE Transactions on Pattern Analysis and Machine Intelligence (TPAMI)*, vol. 38, no. 2, pp. 295–307, 2015.
 - [25] J. Kim, J. K. Lee, and K. M. Lee, "Accurate image super-resolution using very deep convolutional networks," in Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR), 2016.
 - [26] B. Lim, S. Son, H. Kim, S. Nah, and K. M. Lee, "Enhanced deep residual networks for single image super-resolution," in Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition Workshops (CVPRW), 2017.
 - [27] Y. Zhang, K. Li, K. Li, L. Wang, B. Zhong, and Y. Fu, "Image super-resolution using very deep residual channel attention networks," in Proceedings of the European Conference on Computer Vision (ECCV), 2018.
 - [28] Z. Hui, X. Gao, Y. Yang, and X. Wang, "Lightweight image super-resolution with information multi-distillation network," in Proceedings of the ACM International Conference on Multimedia (MM), 2019.
 - [29] C. Ledig et al., "Photo-realistic single image super-resolution using a generative adversarial network," in Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR), 2017.
 - [30] J. Chen et al., "Run, don't walk: Chasing higher FLOPS for faster neural networks," in Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR), 2023.
 - [31] N. Ma, X. Zhang, and J. Sun, "Funnel activation for visual recognition," in Proceedings of the European Conference on Computer Vision (ECCV), 2020.
 - [32] J. Hu, L. Shen, and G. Sun, "Squeeze-and-excitation networks," in Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR), 2018.
 - [33] C. Tian, Y. Xu, W. Zuo, B. Du, C.-W. Lin, and D. Zhang, "Lightweight image super-resolution with enhanced CNN," *Knowledge-Based Syst.*, vol. 205, 106235, 2020.
 - [34] F. Kong, M. Li, S. Liu, D. Liu, J. He, Y. Bai, F. Chen, and L. Fu, "Residual local feature network for efficient super-resolution," in Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition Workshops (CVPRW), 2022.
 - [35] Z. Li, Y. Liu, W. Chen, H. Jiao, and X. Wang, "Blueprint separable residual network for efficient image super-resolution," in Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition Workshops (CVPRW), 2022.
 - [36] J. Liu, J. Tang, and G. Wu, "Residual feature distillation network for lightweight image super-resolution," in Proceedings of the European Conference on Computer Vision Workshops (ECCVW), 2020.