
A REVIEW ON IMAGE RESOLUTION THROUGH SUPER RESOLUTION

Mohammad Rafi¹, Sinchana P², Umme Zaiba³

^{1,2,3}Department of Computer Science and Engineering, UBDTCE, India

ABSTRACT

This survey delves into the realm of image super-resolution (SR), a critical aspect for high-quality displays and medical imaging. Focusing on deep learning algorithms for single-image super-resolution (SISR), the survey categorizes and reviews recent SR techniques, analyzes performance-affecting factors, evaluates algorithms on the Urban100 dataset, and discusses challenges and future directions in the field. Built upon an examination of 12 existing review papers, the survey sheds light on the complexities of this ill-posed inverse problem. The study introduces innovative SR approaches, including SR3 using denoising diffusion, MSRN with multi-scale residual blocks, SAN with second-order channel attention, LapSRN employing a Laplacian Pyramid Network, and CinCGAN for unsupervised learning.

Further contributions include the RCAN with residual channel attention, IDN with information distillation, IMDN balancing performance and applicability, RDN leveraging hierarchical features, SRFBN employing a feedback network, and HAN focusing on texture detail preservation. These approaches showcase advancements in SR methodologies, addressing challenges and providing efficient solutions for accurate image super-resolution. Experimental evaluations and comparisons demonstrate the effectiveness and superiority of these models, paving the way for future developments in the field.

1. INTRODUCTION

Image super-resolution (SR) has become a hot topic in recent years. SR involves turning a blurry image into a sharp one, crucial for things like high-quality displays and medical imaging. It's a tricky problem, as there can be multiple solutions to the same low-resolution image, making it what experts call an ill-posed inverse problem. This complexity increases with higher scaling factors, leading to potential errors in reproduced information.

There are two main types of SR methods: traditional and deep learning. While traditional methods have been around for a long time, deep learning, a branch of machine learning, has shown better results in recent years. This survey focuses on deep learning algorithms for single image super-resolution (SISR).

The authors of the survey make five main contributions:

- 1) a thorough review of recent SR techniques,
- 2) a new way to categorize SR algorithms based on their structures,
- 3) a comprehensive analysis of various factors affecting performance,
- 4) a systematic evaluation of algorithms on Urban100 dataset, and
- 5) discussing challenges and potential future directions in the field.

In addition to the aforementioned insights, it is noteworthy that our comprehensive survey is built upon an extensive examination of 12 existing review papers on Image Super-Resolution (ISR). The SR3 introduces an innovative single-image super-resolution approach inspired by Denoising Diffusion Probabilistic Models. Using a U-Net architecture with a denoising objective, SR3 iteratively removes noise from images, minimizing a defined loss function and maintaining a constant inference step count. It demonstrates effectiveness across magnification factors, allowing efficient cascading for training and inference. Human evaluations are employed for accurate quality assessment, recognizing limitations of automated metrics like PSNR and SSIM in specific scenarios [1].

The Multi-scale Residual Network (MSRN) addresses drawbacks in Single-Image Super-Resolution (SISR) models. It introduces a novel multi-scale residual block (MSRB) for adaptive feature detection and fusion at different scales. MSRN's simplicity allows easy adaptation to various upscaling factors without requiring complex training tricks. The model surpasses state-of-the-art methods in SISR without deep network structures and demonstrates versatility in computer vision tasks. Its contributions include the innovative MSRB, efficient hierarchical features fusion, and superior performance without relying on intricate training methodologies [2].

The proposed Second-Order Attention Network (SAN) enhances Single Image Super-Resolution (SISR) by introducing a Second-Order Channel Attention (SOCA) mechanism for adaptive feature learning. A Non-Locally Enhanced Residual Group (NLRG) captures long-distance spatial contextual information, improving visual quality. SAN outperforms existing methods in both quantitative and visual aspects, showcasing its superiority in accurate image SR. The SOCA mechanism focuses on informative features, and NLRG incorporates non-local operations, addressing limitations in existing CNN-based SR models. The SAN contributions include a deep architecture, the innovative SOCA mechanism, and the NLRG structure for effective feature correlation learning [3].

The LapSRN introduces a deep Laplacian Pyramid Network for Single Image Super-Resolution (SISR), addressing issues in existing methods. It progressively reconstructs high-resolution images in a coarse-to-fine manner, optimizing convolutional layers and upsampling filters jointly. The model achieves both speed and accuracy by avoiding pre-defined upsampling, utilizing deep supervision, and applying Charbonnier loss functions. LapSRN surpasses other CNN-based SISR methods, demonstrating improved mapping complexity and reduced spatial aliasing artifacts. It introduces flexibility in progressive reconstruction and multi-scale training for resource-aware adaptability [4].

The study addresses challenges in Super-Resolution (SR) by proposing the Cycle-in-Cycle Generative Adversarial Networks (CinCGAN) for unsupervised learning. CinCGAN utilizes two CycleGANs, with the first denoising and deblurring the low-resolution input, and the second further upsampling for competitive performance. This approach accommodates scenarios where high-resolution data, downscaling methods, and degradation functions are unknown. Contributions include tackling a general SR problem, exploring unsupervised training, and introducing the stable CinCGAN structure, providing results comparable to supervised CNN networks [5].

The study focuses on single image super-resolution (SR), highlighting challenges in accurately reconstructing high-resolution (HR) images from low-resolution (LR) inputs. It introduces Residual Channel Attention Networks (RCAN), a very deep trainable network with a novel structure (RIR) and channel attention (CA) mechanism. RCAN enhances discriminative learning by adaptively rescaling channel-wise features, surpassing state-of-the-art methods in visual SR results. The contributions include the introduction of RCAN, RIR structure, and CA mechanism for precise image SR [6].

The study addresses single image super-resolution (SISR) challenges by proposing an Information Distillation Network (IDN) with lightweight parameters. IDN employs a feature extraction block, multiple information distillation blocks (DBlocks), and a reconstruction block to progressively distill and aggregate residual information, achieving competitive results with fewer convolutional layers. Unlike deep networks with large computational costs, IDN maintains real-time speed and better reconstruction accuracy. The key component, DBlock, features an enhancement unit for improving representation power and a compression unit to distill useful information. Overall, IDN offers a concise structure for efficient and accurate image super-resolution [7].

The study introduces a Lightweight Information Multi-distillation Network (IMDN) for single image super-resolution (SISR) that balances performance and applicability. Employing an Information Multi-distillation Block (IMDB) inspired by previous work, IMDN achieves competitive results with a modest number of parameters. The proposed Adaptive Cropping Strategy (ACS) allows the network to handle images of arbitrary sizes efficiently. The contributions include a lightweight IMDN with a contrast-aware attention layer, ACS for processing images of any size, and insights into factors affecting actual inference time for guiding lightweight network design [8].

The study addresses Single Image Super-Resolution (SISR), proposing a Residual Dense Network (RDN) to leverage hierarchical features from the original low-resolution (LR) image. Introducing the Residual Dense Block (RDB) as the building module, RDN supports contiguous memory among RDBs, facilitating the efficient extraction of local dense features. Global Feature Fusion (GFF) is employed to adaptively preserve hierarchical features in a global manner. The contributions include the unified RDN framework for diverse degradation models, the innovative RDB with local and global feature connections, and the achievement of high-quality image SR by leveraging hierarchical information from the LR image [9].

The Feedback Network for Image Super-Resolution addresses the ill-posed nature of the task by proposing a Super-Resolution Feedback Network (SRFBN). Leveraging deep learning and feedback connections inspired by cognitive theory, the SRFBN employs a recurrent structure with a feedback block (FB) to refine low-level information using high-level details. The FB, constructed with up- and down-sampling layers and dense skip connections, facilitates top-down feedback flows. The proposed curriculum-based training strategy enhances the SRFBN's ability to handle complex degradation models by arranging target HR images from easy to hard based on recovery difficulty. Experimental results demonstrate the superiority of SRFBN over existing methods [10].

The Holistic Attention Network (HAN) addresses challenges in Single Image Super-Resolution (SISR), focusing on texture detail preservation. Leveraging convolutional neural networks (CNNs), it introduces a Layer Attention Module (LAM) and a Channel-Spatial Attention Module (CSAM) to enhance feature expression and correlation learning. Unlike existing methods, HAN explores correlations among hierarchical layers, channels, and positions of each channel, preventing loss of informative textures. LAM considers multi-scale layer correlations, while CSAM improves discrimination ability, collaboratively enhancing SR results. Extensive experiments demonstrate the effectiveness of HAN against state-of-the-art SISR approaches [11].

2. LITERATURE SURVEY

This literature survey delves into the domain of super-resolution, focusing on the significance of high-resolution images in various applications such as computer displays, HD television sets, and handheld devices. Super-resolution finds applications in object detection, face recognition, medical imaging, remote sensing, astronomical images, and forensics.

The survey identifies super-resolution as a challenging and open research problem in computer vision due to its ill-posed nature and the existence of multiple solutions for the same low-resolution image. The need for reliable prior information and the increased complexity at higher up-scaling factors contribute to the difficulty of the problem. Additionally, assessing the quality of output poses challenges as quantitative metrics like PSNR and SSIM may not directly correlate with human perception.

The study classifies super-resolution methods into traditional and deep learning categories. While classical algorithms have been present for decades, recent advancements have shown that deep learning-based methods outperform them. Deep learning, a branch of machine learning, is employed to automatically learn relationships between input and output directly from data. The survey primarily focuses on deep learning algorithms for super-resolution, given their promising results in various AI sub-fields.

The literature then narrows down its focus to Single Image Super-Resolution (SISR) and categorizes existing methods into nine groups based on distinctive features in their model designs.

The taxonomy as shown in Figure 1 encompasses linear networks, among other designs, highlighting the evolution of SISR techniques. The survey emphasizes the importance of high-resolution images, outlines the challenges in super-resolution, and underscores the dominance of deep learning over traditional methods. It provides a detailed taxonomy for existing SISR techniques, showcasing the progression of model designs in this specific sub-field [12].



Fig. 1. The taxonomy of the existing single-image super-resolution techniques based on the most distinguishing features [12].

Datasets:

The study evaluates state-of-the-art super-resolution algorithms on benchmark datasets: Set5, Set14, BSD100, Urban100, DIV2K, and Manga109.

- **Set5** is a small classical dataset with five test images,
- **Set14** has 14 test images.
- **BSD100** includes 100 diverse test images,
- **Urban100** focuses on urban scenes, and
- **DIV2K**, used for NITRE challenge, comprises 800 training and 100 each for testing and validation at 2K resolution.
- **Manga109**, a recent addition, consists of 109 manga volume images drawn by Japanese artists.

The study provides visual comparisons of algorithms as shown in Figure 2 aiming to improve PSNR and GAN-based algorithms focusing on perceptual enhancement. While GAN-based outputs appear more crisp, their PSNR values are lower compared to pixel-level loss methods.

8× Super-resolution:

The study notes that most algorithms struggle to reproduce textures in high-magnification versions of images, as evident in comparisons for 8× super-resolution.(in Table 1 and Figure 2 the comparisons are provided for 8× super-resolution)

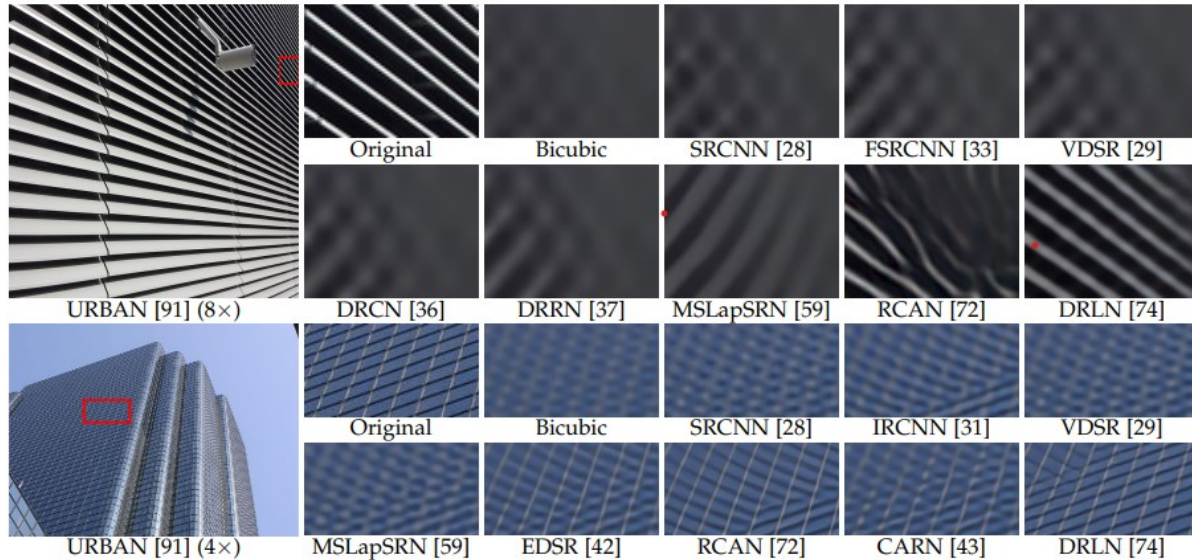


Fig. 2. Super-resolution comparison on 8× and 4× sample images with sharp edges and texture, taken from URBAN100[91].

TABLE 1- The performance of state-of-the-art algorithms on widely used publicly available datasets, in terms of PSNR (in dB) and SSIM for 8×.

Scale	Method	SET5		SET14		BSD100		URBAN100		MANGA109	
		PSNR	SSIM	PSNR	SSIM	PSNR	SSIM	PSNR	SSIM	PSNR	SSIM
8×	Bicubic	24.40	0.6580	23.10	0.5660	23.67	0.5480	20.74	0.5160	21.47	0.6500
	SRCNN	25.33	0.6900	23.76	0.5910	24.13	0.5660	21.29	0.5440	22.46	0.6950
	FSRCNN	20.13	0.5520	19.75	0.4820	24.21	0.5680	21.32	0.5380	22.39	0.6730
	SCN	25.59	0.7071	24.02	0.6028	24.30	0.5698	21.52	0.5571	22.68	0.6963
	VDSR	25.93	0.7240	24.26	0.6140	24.49	0.5830	21.70	0.5710	23.16	0.7250
	LapSRN	26.15	0.7380	24.35	0.6200	24.54	0.5860	21.81	0.5810	23.39	0.7350
	MemNet	26.16	0.7414								
	MSLapSRN	26.34	0.7558	24.38	0.6199	24.58	0.5842	21.89	0.5825	23.56	0.7387
				24.57	0.6273	24.65	0.5895	22.06	0.5963	23.90	0.7564
	EDSR	26.96	0.7762	24.91	0.6420	24.81	0.5985	22.51	0.6221	24.69	0.7841
	D-DBPN	27.21	0.7840	25.13	0.6480	24.88	0.6010	22.73	0.6312	25.14	0.7987
	RCAN	27.31	0.7878	25.23	0.6511	24.98	0.6058	23.00	0.6452	25.24	0.8029

The table presents a comparative analysis of super-resolution methods, demonstrating that advanced deep learning models like EDSR, D-DBPN, RCAN, DRLN, and EBRN consistently outperform traditional methods across various datasets, showcasing their efficacy in achieving superior image quality in terms of PSNR and SSIM metrics at different scales.

- **PSNR** (Peak Signal-to-Noise Ratio) measures image fidelity by evaluating the ratio of maximum signal strength to noise.
- **SSIM** (Structural Similarity Index) assesses perceived image quality based on structural information and luminance.

Super-Resolution Methods Overview:

1. **Bicubic**: Bicubic Interpolation, a mathematical algorithm used for high-quality image resizing.
2. **SRCNN**: Super-Resolution Convolutional Neural Network, a deep learning model for single-image super-resolution.
3. **FSRCNN**: Fast Super-Resolution Convolutional Neural Network, an accelerated version of SRCNN.
4. **SCN**: Sparse Coding Network, a model utilizing sparse representations for image super-resolution.
5. **VDSR**: Very Deep Super-Resolution, a deep neural network designed for high-performance single-image super-resolution.
6. **LapSRN**: Laplacian Pyramid Super-Resolution Network, a model using a laplacian pyramid to achieve super-resolution.
7. **MemNet**: Memory Network, a neural network architecture incorporating memory blocks for image super-resolution.
8. **MSLapSRN**: Multi-Scale Laplacian Pyramid Super-Resolution Network, a model combining multi-scale and laplacian pyramid approaches for super-resolution.
9. **EDSR**: Enhanced Deep Super-Resolution, a deep learning model with enhanced architecture for improved super-resolution.
10. **D-DBPN**: Deep Back-Projection Networks, a deep learning model designed to efficiently handle super-resolution tasks.
11. **RCAN**: Residual Channel Attention Networks, a network incorporating attention mechanisms for better super-resolution performance.
12. **DRLN**: Deep Recursive Residual Learning for Single Image Super-Resolution, a model utilizing recursive learning for improved super-resolution.
13. **EBRN**: Enhanced Back-Projection Networks, a model with enhanced back-projection techniques for better super-resolution results.

3. METHODOLOGY

Methodological Insights into Image Super-Resolution: Unveiling Advancements and Output Results.

1. The survey reviews 11 existing papers on image super-resolution methodologies.
2. Each paper implements specific algorithms and approaches for super-resolution tasks.
3. Methodologies include SR3, MSRN, SAN, LapSRN, CinCGAN, RCAN, IDN, IMDN, RDN, SRFBN, and HAN.
4. The survey categorizes these methodologies based on their structural characteristics.
5. Performance factors affecting these methodologies are systematically analyzed.
6. Evaluation is conducted on the Urban100 dataset to assess algorithm effectiveness.
7. The study introduces a new categorization approach based on algorithm structures.
8. Comprehensive analyses of factors influencing super-resolution performance are provided.
9. Experimental evaluations include assessing image output results of the implemented methodologies.
10. The survey discusses challenges and outlines potential future directions in the field of image super-resolution.

Image Super-Resolution via Iterative Refinement [1]: SR3 models exhibit effectiveness across various magnification factors and input resolutions, demonstrating their adaptability (see Figure 3). The models can be cascaded, transitioning from 64x64 to 256x256 and then to 1024x1024, offering flexibility in training smaller models independently rather than a single large model with high magnification. Cascading enhances efficiency in inference, as it allows for more cost-effective iterative refinement steps, particularly in generating high-resolution images. Additionally, experiments showcase the feasibility of cascading an unconditional generative model with SR3, enabling the unconditional generation of high-fidelity images. The experiments encompass both natural and face image domains.

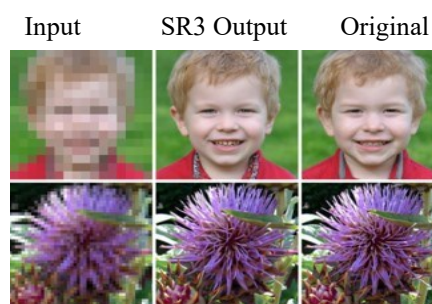


Fig. 3. Two representative SR3 outputs: (top) 8× face super-resolution at 16×16→128×128 pixels (bottom) 4× natural image super-resolution at 64× 64→256× 256 pixels.

Evaluation

The SR3 models undergo evaluation in two domains:

1. Face Super-Resolution: Evaluated at 16x16, 128x128, and 64x64 resolutions, trained on FFHQ, and tested on CelebA-HQ.
2. Natural Image Super-Resolution: Assessed at 64x64 to 256x256 and 56x56 to 224x224 resolutions using ImageNet.
3. Additionally, the evaluation includes:
 - Unconditional 1024x1024 face generation using a cascade of 3 SR3 models.
 - Class-conditional 256x256 ImageNet image generation via a cascade of 2 SR3 models.

SR3 is compared against EnhanceNet, ESRGAN, SRFlow, FSRGAN, PULSE, and a Regression baseline with matching architecture and model capacity. This comparative analysis, spanning qualitative and quantitative assessments, involves human evaluation, FID scores, and classification accuracy using a pre-trained model on super-resolution outputs. Importantly, the Regression baseline allows a direct examination of the benefits of iterative refinement compared to a single-step regression model, eliminating the influences of model size, architecture, and training data.

Cascaded Image Generation with SR3 Models

The study investigates cascaded image generation as shown in Figure 4, where SR3 models at different scales are combined with generative models to facilitate high-resolution image synthesis. Cascading enables parallel training of models, each addressing simpler tasks with fewer parameters and less computation. This approach proves more efficient for iterative refinement models, particularly when using more refinement steps at low resolutions and fewer steps at higher resolutions. In cascaded face generation, a DDPM model for unconditional 64x64 face images is trained, followed by two 4x SR3 models, ultimately reaching resolutions of 256x256 and 1024x1024 pixels. Data augmentation, involving Gaussian Blurring noise during SR3 training, significantly improves robustness and FID scores. Initial cascade experiments yield lower FID scores than models trained directly at full resolution, emphasizing efficiency. The study also explores the denoising objective's Lp norm, finding that L1 norm provides slightly better FID scores, while subsequent work indicates that L2 norm generates greater diversity in SR3 outputs [1]. SR3 employs conditional diffusion models for single-image super-resolution, excelling on natural and face images across various magnification factors. Human studies reveal fool rates close to 50% on faces and 40% on natural images, indicating high-fidelity outputs. Despite concerns about biases and computation costs, SR3's effectiveness suggests its utility in reducing dataset bias by generating synthetic data for underrepresented groups.



Fig. 4. Cascaded generation with an unconditional model chained with two SR3 models.

Multi-scale Residual Network for Image Super-Resolution [2]:

The methodology follows a three-fold augmentation of training data, involving scaling, rotation, and flipping. Each training batch randomly extracts 16 Low-Resolution (LR) patches sized 64x64, with an epoch consisting of 1000 back-propagation iterations. The model employs the ADAM optimizer with a learning rate (lr) set to 0.0001. The final model incorporates 8 Multi-Scale Residual Blocks (MSRB) with 64 feature maps in each MSRB output and bottleneck layer. The MSRN is implemented using the PyTorch framework and trained on an NVIDIA Titan Xp GPU. No special weight initialization or training tricks are used, and the code is available at <https://github.com/MIVRC/MSRN-PyTorch>. In Qualitative Analysis, the authors highlight the benefits of the proposed Multi-Scale Residual Block (MSRB). The MSRB is presented as an efficient feature extraction structure capable of adaptively detecting image features at various scales. Comparative experiments are conducted against residual blocks, dense blocks, and MSRB in Single Image Super-Resolution (SISR) tasks. Results demonstrate the superiority of MSRB across all upsampling factors, supported by visualizations indicating sparse activations, where MSRB produces more valid activation maps. The authors explore the impact of increasing the number of MSRBs on network performance in terms of depth. The study reveals that the performance of the MSRN improves with the growing number of MSRBs, although a balance is sought between enhanced performance and increased network complexity. Ultimately, 8 MSRBs are chosen, delivering results close to EDSR while maintaining a more computationally efficient model.

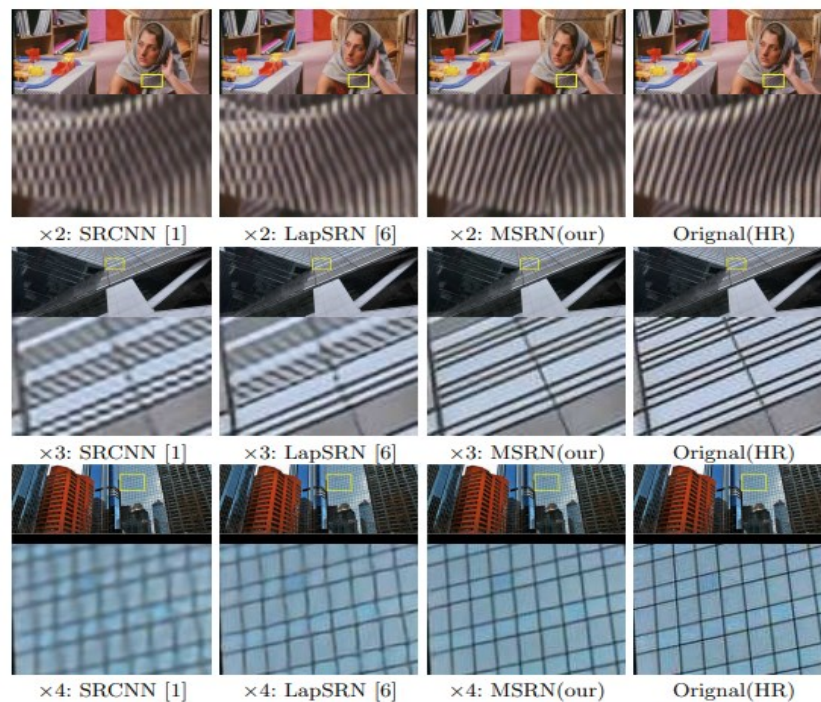


Fig. 5. Visual comparison for $\times 2$, $\times 3$, $\times 4$ SR images. Our MSRN can reconstruct realistic images with sharp edges.

Furthermore, the proposed MSRB module is applied to other low-level computer vision tasks, such as image denoising and image dehazing, demonstrating promising results and validating the module's effectiveness beyond SISR tasks.

The paper introduces an efficient Multi-Scale Residual Block (MSRB) for adaptive feature detection, forming the basis for the Multi-Scale Residual Network (MSRN) as a straightforward and effective Super-Resolution model. The MSRN leverages local multi-scale and hierarchical features, yielding accurate SR images. Additionally, the MSRB module demonstrates promising results in tasks like image denoising and dehazing across computer vision applications.

Second-Order Attention Network for Single Image Super-Resolution [3]:

Experiments:

Setup:

- Utilizes 800 high-resolution images from DIV2K dataset for training.
- Testing involves 5 benchmark datasets: Set5, Set14, BSD100, Urban100, and Manga109.
- Experiments conducted with Bicubic (BI) and Blur-downscale (BD) degradation models.
- Evaluation based on PSNR and SSIM metrics in the Y channel of transformed YCbCr space.
- Training involves data augmentation: random rotations and horizontal flipping.
- Model input consists of 8 LR color patches (48×48) in each mini-batch.
- Implemented on PyTorch framework using an Nvidia 1080Ti GPU.

Ablation Study:

- Two main components in the proposed SAN: Non-locally Enhanced Residual Group (NLRG) and Second-order Channel Attention (SOCA).
- NLRG variants (Ra to Rf) tested on Set5 dataset, showing the effectiveness of individual modules.
- SOCA variants (Rd to Ri) demonstrate the impact of second-order channel attention.
- Results indicate the superiority of the proposed SOCA over first-order channel attention.

Results with Bicubic Degradation (BI):

- Compared SAN with 11 state-of-the-art CNN-based SR methods.
- Self-ensemble method (SAN+) applied to enhance results.
- SAN+ outperforms other methods on all datasets and scaling factors.
- SAN and RCAN show similar performance, surpassing other methods due to their adoption of channel attention.
- SAN excels on datasets with rich texture information, while RCAN performs better on datasets with repeated edge information.
- Visual quality comparison highlights as shown in Figure 6 SAN's ability to reconstruct lattices accurately, recovering high-frequency details and demonstrating superior representational ability.

Summary of Proposed Methodology:

The proposed methodology introduces a deep Second-Order Attention Network (SAN) for accurate image Super-Resolution (SR). The key components include:

1. Non-locally Enhanced Residual Group (NLRG):

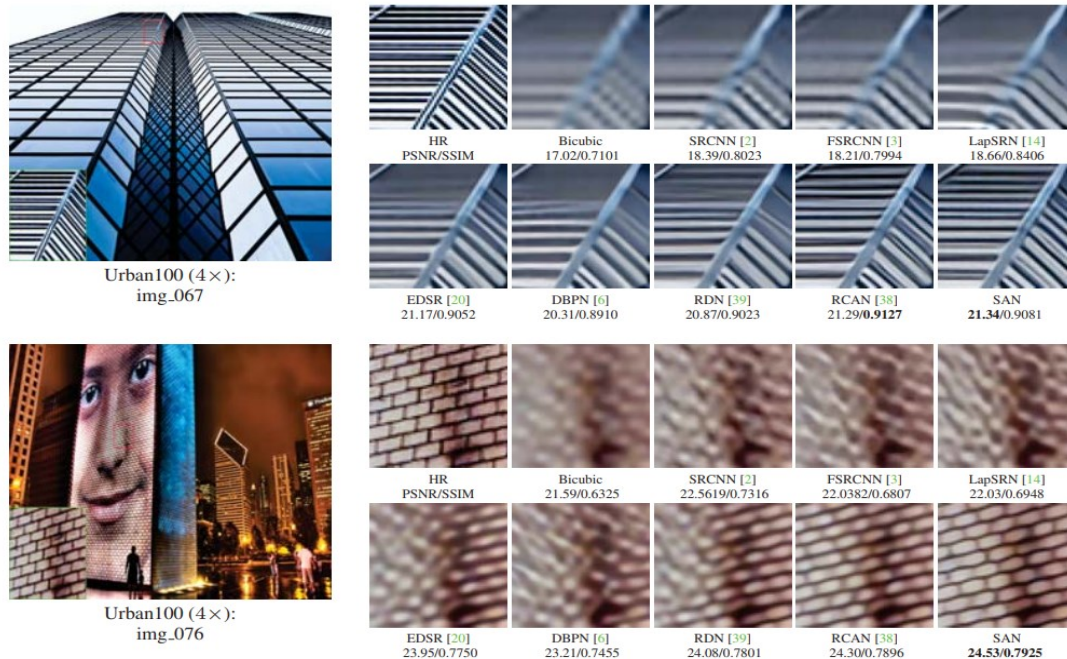


Fig. 6. Visual comparison for 4× SR with BI model on Urban100 dataset. The best results are highlighted

- Structure designed to capture long-distance dependencies and structural information.
 - Incorporates non-local operations within the network.
 - Enables the bypassing of abundant low-frequency information from Low-Resolution (LR) images through share-source skip connections.
- #### 2. Second-Order Channel Attention (SOCA) Module:
- Aims to learn feature interdependencies.
 - Utilizes global covariance pooling to gather discriminative representations.

The methodology focuses on exploiting both spatial feature correlations through NLRG and channel feature correlations through SOCA. Extensive experiments, including Super-Resolution with Bicubic (BI) and Blur-Downscale (BD) degradation models, demonstrate the effectiveness of the proposed SAN. The evaluation is based on quantitative metrics and visual assessments, indicating promising results in terms of both accuracy and visual quality.

Fast and accurate ISR with Deep Laplacian Pyramid Networks [4]:

The various components of the proposed network Model Design for image super-resolution:

Model Design:

- They train a LapSRN model with 5 convolutional layers at each pyramid level.
 - They analyze the impact of pyramid network structure, global residual learning, robust loss functions, and multi-scale supervision.
- #### 1. Pyramid Structure:
- Removing the pyramid structure leads to a network similar to FSRCNN but with global residual learning.
 - Results show that the pyramid structure significantly improves performance, validating the effectiveness of the Laplacian pyramid network design.
- #### 2. Global Residual Learning:
- The authors demonstrate the effectiveness of global residual learning by comparing a non-residual network with the full LapSRN model.
 - The full model outperforms the non-residual network, particularly within a short training period.
- #### 3. Loss Function:
- The effectiveness of the Charbonnier loss function is validated by comparing it with a conventional L2 loss function.
 - Results show that the network optimized with Charbonnier loss achieves comparable performance with SRCNN in fewer iterations.

4. Multi-scale Supervision:

- Multiple loss functions are used to supervise the intermediate output at each pyramid level.
- Results indicate that multi-scale supervision helps in progressively reconstructing high-resolution images and reducing spatial aliasing artifacts.

Employ a multi-scale training strategy for their LapSRN model, focusing on $2^n \times$ samples. They train the LapSRNSS-D5R8 model with various scale combinations: $\{2 \times\}$, $\{4 \times\}$, $\{8 \times\}$, $\{2 \times, 4 \times\}$, $\{2 \times, 8 \times\}$, $\{4 \times, 8 \times\}$, and $\{2 \times, 4 \times, 8 \times\}$, distributing batches equally across different upsampling scales. Despite having the same number of parameters due to parameter sharing, these models are evaluated for $2 \times$, $4 \times$, and $8 \times$ super-resolution, with an additional evaluation for $3 \times$ SR using a 2-level LapSRN. Experimental results indicate that the model trained across multiple scales effectively handles different upsampling scales and generalizes well to unseen $3 \times$ SR examples. Multi-scale models exhibit favorable performance compared to single-scale models, especially on the URBAN100 dataset. Complete quantitative and visual comparisons are available in the supplementary material.

The experimental results and comparisons for their proposed LapSRN model:

1. Objective Comparisons:

- The LapSRN is compared with 10 state-of-the-art SR algorithms, including dictionary-based, self-similarity based, and CNN-based methods.
- Experiments are conducted on five benchmark datasets with natural scenes, urban scenes, and Japanese manga.
- Evaluation metrics include PSNR, SSIM, and IFC for $2 \times$, $3 \times$, $4 \times$, and $8 \times$ SR.

2. Variations of LapSRN:

- Three variations of LapSRN are compared: LapSRNSS-D5R2, LapSRNSS-D5R5, and LapSRNSS-D5R8, each with different depths.
- The multi-scale training strategy with $2 \times$, $4 \times$, and $8 \times$ SR samples is employed, denoted as MS-LapSRN.

3. Quantitative Results:

- LapSRN demonstrates superior performance, especially on $4 \times$ and $8 \times$ SR, with higher IFC values.
- The method achieves comparable results to DRRN without using $3 \times$ SR samples for training.

4. Visual Comparisons (Figure 7):

- Visual comparisons on various datasets for $4 \times$ and $8 \times$ SR demonstrate LapSRN's ability to accurately reconstruct parallel lines, grids, and texts.
 - The method effectively suppresses artifacts caused by spatial aliasing in contrast to pre-upsampling-based methods.
 - For $8 \times$ SR, LapSRN outperforms other methods in reconstructing fine structures at a relatively fast speed.
- #### 5. Additional Analysis:
- Comparison with direct reconstruction-based methods, like VDSR, is provided in the supplementary material, exploring the progressive reconstruction strategy.

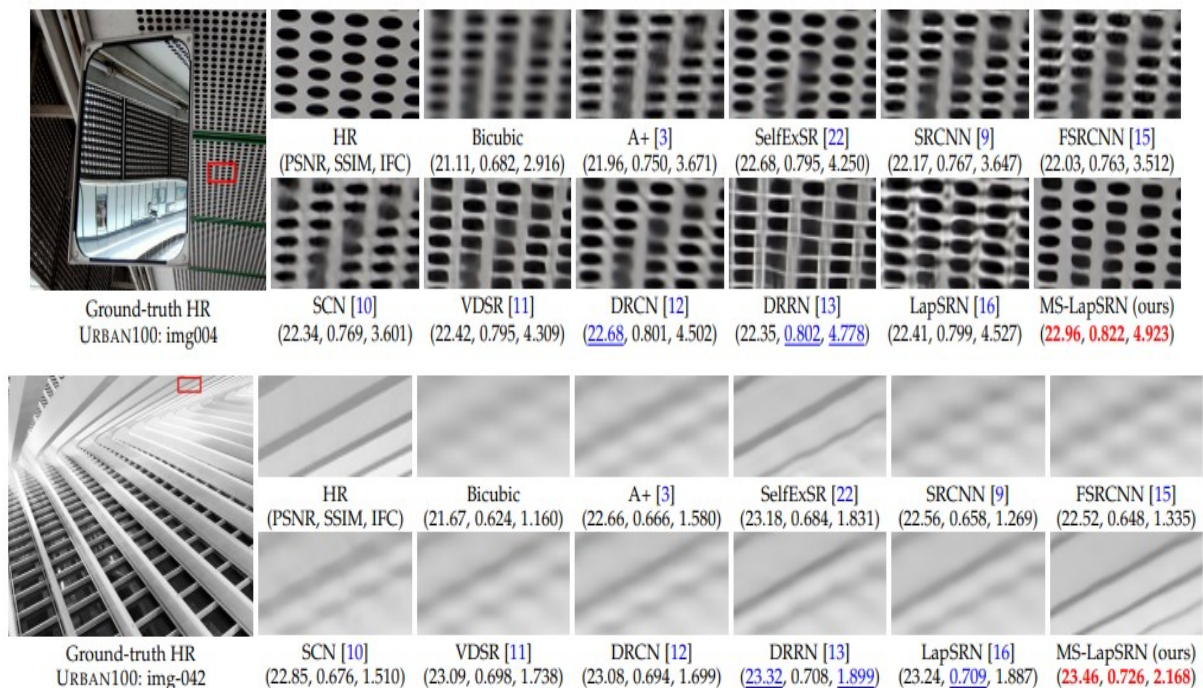


Fig. 7. Visual comparison for $8 \times$ SR on the URBAN100 datasets.

6. Availability and Human Subject Study:

- The authors provide their source code and SR results for all evaluated methods on their project website.
- A human subject study using pairwise comparison is conducted, with detailed analysis available in the supplementary material.

7. Limitations:

- The section concludes with a discussion of the proposed method's limitations.

The experimental results and comparisons demonstrate the effectiveness of LapSRN, especially in achieving high-quality super-resolution, overcoming spatial aliasing artifacts, and outperforming other state-of-the-art methods.

We propose a novel LapSRN model for image super-resolution, employing a deep convolutional network in a Laplacian pyramid framework. Our approach achieves superior performance with 73% fewer parameters than the preliminary method, leveraging multiscale training and local skip connections for efficient and accurate results applicable to diverse image transformation tasks.

Unsupervised Image Super-Resolution using Cycle-in-Cycle Generative Adversarial Networks [5]:

The proposed method addresses the single image super-resolution problem by introducing a solution that involves jointly fine-tuning low-resolution (LR) to high-resolution (HR) networks using CinCGAN. The process includes sequential updates to LR→clean LR and LR→HR models, with constraints L_{LR_total} and L_{HR_total} . The G1 network denoises and deblurs the degraded input image, while the SR network up-samples and further restores the intermediate image.

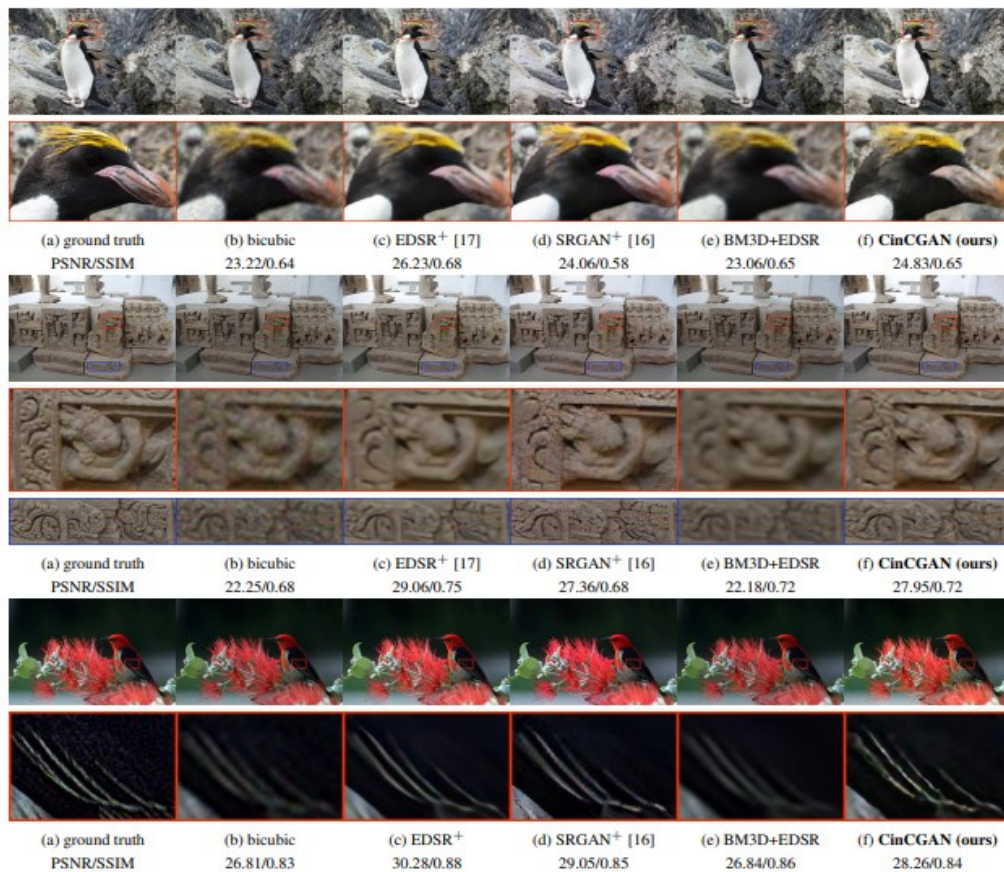


Fig. 8. Super-resolution results of “0801”, “0816” and “0853” (DIV2K) with scale factor $\times 4$. EDSR+ and SRGAN+ are trained on paired NTIRE2018 track 2 dataset. BM3D+EDSR means using BM3D for denoising first and then using EDSR for super-resolution. The proposed CinCGAN model shows comparable results with SRGAN+ and is better than BM3D+EDSR method.

Table 2- Quantitative evaluation on NTIRE 2018 track 2 dataset of the proposed CinCGAN model, in terms of PSNR and SSIM.

method	bicubic	FSRCNN [4]	EDSR [17]	EDSR+	SRGAN+ [16]	BM3D+EDSR	CinCGAN (ours)
PSNR	22.85	22.79	22.67	25.77	24.33	22.88	24.33
SSIM	0.65	0.61	0.62	0.71	0.67	0.68	0.69

The final super-resolved (SR) image demonstrates superior visual results compared to alternative structures. The approach employs unsupervised learning, inspired by image-to-image translation applications, utilizing generative adversarial networks (GANs). The method involves two CycleGANs, with the second GAN covering the first, and a three-step pipeline: mapping input LR images to clean and bicubic-downsampled LR space, up-sampling the intermediate result with a deep model, and fine-tuning the modules in an end-to-end manner. Experimental results show that the proposed unsupervised method achieves comparable performance to state-of-the-art supervised models.

Image Super-Resolution Using Very Deep Residual Channel Attention Networks [6]:

The study introduces Very Deep Residual Channel Attention Networks (RCAN) as a solution for accurate image super-resolution (SR). Visual comparisons reveal that RCAN surpasses existing methods, particularly excelling in handling challenging images with intricate details, mitigating blurring artifacts, and recovering more informative components. The proposed RCAN leverages a residual in residual (RIR) structure to achieve significant depth, allowing for the effective learning of both low and high-frequency information. Additionally, the incorporation of a channel attention (CA) mechanism enhances adaptability by rescaling channel-wise features, considering interdependencies among channels. The effectiveness of RCAN is further demonstrated through extensive experiments on SR with both bicubic (BI) and realistic (BD) degradation models, showcasing promising results not only in image enhancement but also in object recognition tasks.

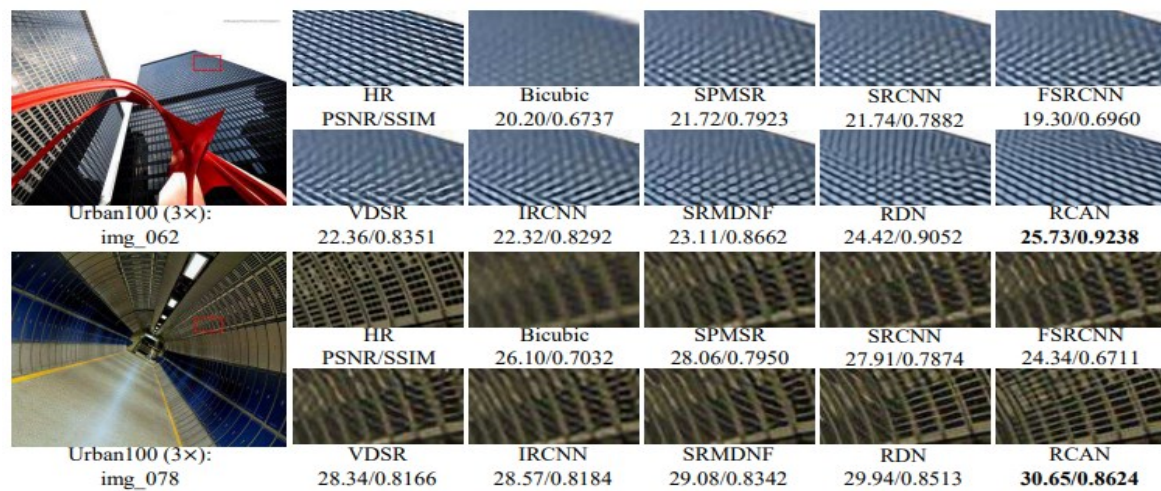


Fig. 9. Visual comparison for 3× SR with BD model on Urban100 dataset. The best results are highlighted

Furthermore, the evaluation of RCAN's object recognition performance reinforces its efficacy. Using ResNet-50 and the ImageNet CLS-LOC dataset, the study compares RCAN with four state-of-the-art methods, including DRCN, FSRCNN, PSyCo, and ENet-E. RCAN consistently achieves the lowest top-1 and top-5 errors, underscoring its robust representational ability. The results highlight RCAN's versatility and its potential as a powerful tool not only for image super-resolution but also as a beneficial pre-processing step for high-level visual tasks like object recognition.

Fast and Accurate Single Image Super-Resolution via Information Distillation Network [7]:

The paper introduces a novel approach called Information Distillation Network (IDN) for Single Image Super-Resolution

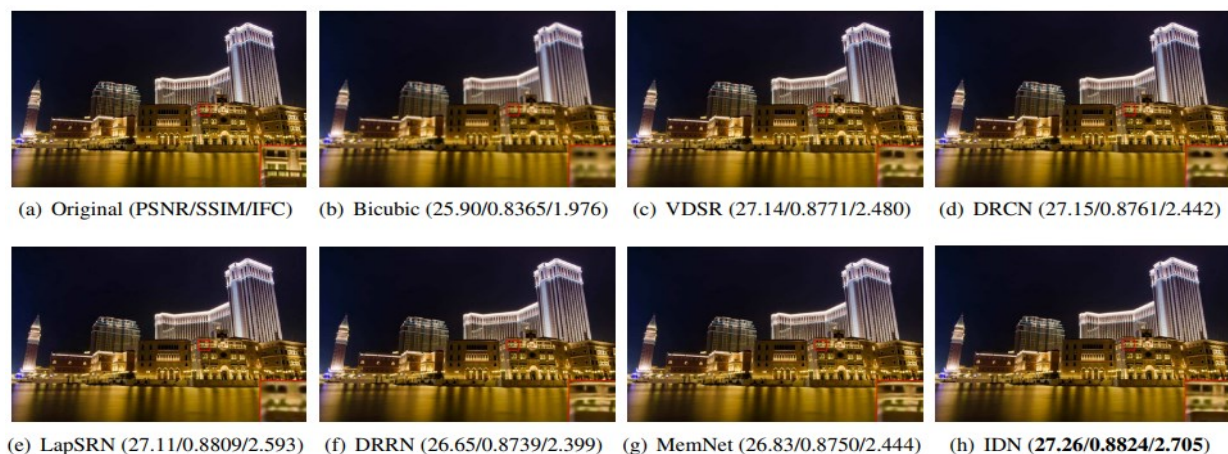


Fig. 10. The “img085” image from the Urban100 dataset with an upscaling factor 4.

(SR). The proposed method is compared with various SR methods on benchmark datasets using peak signal-to-noise ratio (PSNR), structural similarity (SSIM), and information fidelity criterion (IFC) metrics. The results indicate that IDN performs favorably against state-of-the-art methods, particularly outperforming MemNet by a considerable margin. Visual comparisons in Figure 10 demonstrate the effectiveness of IDN in recovering high-frequency information and producing clearer contours without serious artifacts. The proposed method excels in scenarios where other methods introduce fake information or struggle with image quality.

The paper acknowledges a performance difference on the Urban100 dataset and for $3\times$, $4\times$ scale factors, attributing it to MemNet's use of an interpolated low-resolution (LR) image as input, providing more information to the network. IDN, on the other hand, predicts more pixels from scratch, especially in larger images and magnification factors.

In terms of inference time, IDN is reported to be approximately 500 times faster than MemNet with a $2\times$ magnification on the Urban100 dataset. The conclusion highlights IDN's competitive results in terms of PSNR, SSIM, and IFC, as well as its significantly improved inference time compared to state-of-the-art methods.

The authors suggest potential applications in image restoration beyond super-resolution, such as denoising and compression artifacts reduction.

This introduces a novel network using distillation blocks for high-resolution image reconstruction, demonstrating competitive performance on benchmark datasets in terms of PSNR, SSIM, and IFC. Notably, the proposed method achieves significantly faster inference times compared to state-of-the-art methods like DRRN and MemNet. The compact nature of this network suggests broad practical applicability, and future exploration may extend its use to address various image restoration challenges, including denoising and compression artifacts reduction.

Lightweight Image Super-Resolution with Information Multi-distillation Network [8]:

4. EXPERIMENTS

Datasets and Metrics:

- Training Dataset: Utilized the DIV2K dataset with 800 high-quality RGB training images, commonly used in image restoration tasks.
- Evaluation Datasets: Employed benchmark dataset Urban100.
- Metrics: Assessed the performance using peak signal-to-noise ratio (PSNR) and structure similarity index (SSIM). Calculated these metrics on the luminance channel (Y channel of YCbCr channels converted from RGB channels).
- Additional Dataset for Unknown Scale Factor Experiments: Used the RealSR dataset from NTIRE2019 Real Super-Resolution Challenge, containing real low and high-resolution paired images.

Implementation Details:

- LR DIV2K Training Images: Obtained low-resolution (LR) DIV2K training images by downscaling high-resolution (HR) images with scaling factors ($\times 2$, $\times 3$, and $\times 4$) using bicubic interpolation in MATLAB R2017a.
- Input Image Patches: HR image patches with a size of 192×192 were randomly cropped from HR images to serve as input for the model.
- Mini-Batch Size: Set the mini-batch size to 16.
- Data Augmentation: Applied random horizontal flip and 90-degree rotation for data augmentation.
- Training Details: Utilized the ADAM optimizer with a momentum parameter ($\beta_1 = 0.9$). The initial learning rate was set to 2×10^{-4} and halved every 2×10^5 iterations.
- Number of IMDB: Set the number of Information Multi-distillation Blocks (IMDB) to 6 in both IMDN and IMDN_AS.
- Implementation Framework: Used the PyTorch framework for implementing the proposed network.
- Hardware Specifications: The model was trained on a desktop computer equipped with a 4.2GHz Intel i7-7700K CPU, 64GB RAM, and an NVIDIA TITAN Xp GPU with 12GB memory.

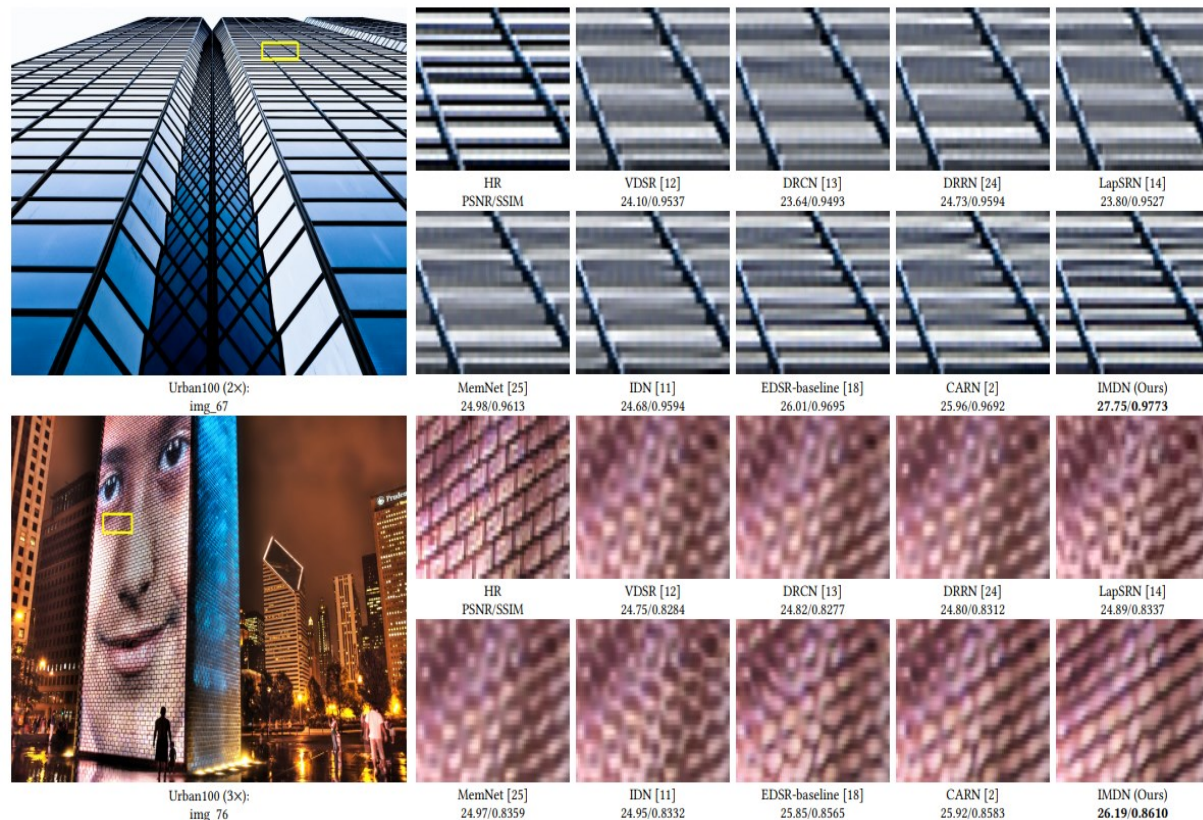


Fig. 11. Visual comparisons of IMDN with other SR methods on Set5 and Urban100 datasets.

In the comparison with 11 state-of-the-art methods, including SRCNN, FSRCNN, VDSR, DRCN, LapSRN, DRRN, MemNet, IDN, EDSR-baseline, SRMDNF, and CARN, the Information Multi-distillation Network (IMDN) demonstrated superior performance. Quantitative assessments for scaling factors $\times 2$, $\times 3$, and $\times 4$ in super-resolution (SR) revealed that IMDN outperformed the other methods across various datasets, particularly excelling at a scaling factor of $\times 2$.

Visual comparisons presented in Figure 11 for scaling factors $\times 2$, $\times 3$ on Urban100 datasets further support these findings. Notably, for the "img_67" image from the Urban100 dataset, IMDN exhibited superior recovery of grid structures compared to other methods. This visual evidence underscores the effectiveness of IMDN in image super-resolution, reinforcing its competitive edge over state-of-the-art approaches.

The Information Multi-distillation Network designed for achieving lightweight and accurate single-image super-resolution. The key innovation lies in a progressive refinement module that systematically extracts hierarchical features. This module collaborates with a contrast-aware channel attention module, leading to a significant and consistent enhancement in super-resolution (SR) performance. Additionally, we propose an adaptive cropping strategy to address the SR challenges associated with arbitrary scale factors, crucial for real-world applications. Through numerous experiments, our method demonstrates a commendable balance among practical considerations, including visual quality, execution speed, and memory consumption. Looking ahead, we envision extending this approach to facilitate other image restoration tasks such as denoising and enhancement.

Residual Dense Network for Image Super-Resolution [9]:

In the Residual Dense Network (RDN) for Image Super-Resolution, the authors provide details about the settings, datasets, metrics, degradation models, and training parameters used in their study. They evaluate the performance of their proposed RDN on various benchmark datasets and degradation models, comparing it with state-of-the-art super-resolution methods.

5. EXPERIMENTAL RESULTS

Settings

Datasets and Metrics:

- The DIV2K dataset (2K resolution) by Timofte et al. is used for training and testing.
- Five benchmark datasets (Set5, Set14, B100, Urban100, Manga109) are employed for testing.
- Evaluation metrics include PSNR and SSIM on the luminance channel (Y) in the transformed YCbCr space.

Degradation Models:

- Three degradation models are used to simulate low-resolution (LR) images: BI (bicubic downsampling), BD (blur + downsampling), and DN (bicubic downsampling + Gaussian noise).

Training Setting:

- Training involves randomly extracting 16 LR RGB patches (32x32 size) per batch.
- Data augmentation includes horizontal/vertical flipping and 90-degree rotation.
- The RDN is implemented using the Torch7 framework with the Adam optimizer.
- Learning rate starts at 10^{-4} for all layers and decreases by half every 200 epochs.
- Training takes approximately 1 day with a Titan Xp GPU for 200 epochs.

Results with BI Degradation Model

- RDN is compared with several state-of-the-art image super-resolution methods using the BI degradation model (bicubic downsampling).
- Results show RDN performs the best on all datasets for scaling factors $\times 2$, $\times 3$, and $\times 4$.
- RDN outperforms other models, including SRDenseNet and MemNet, indicating the effectiveness of the residual dense block (RDB) over alternative architectures.

Results with BD and DN Degradation Models

- RDN is compared with other methods using BD (blur + downsampling) and DN (bicubic downsampling + Gaussian noise) degradation models.
- RDN consistently outperforms other methods in terms of PSNR and SSIM on various datasets.
- Visual comparisons demonstrate RDN's ability to suppress blurring artifacts (BD) and efficiently handle noise while recovering more details (DN).
- The results indicate that RDN is applicable for joint image denoising and super-resolution.

The experimental results showcase the superior performance of the proposed Residual Dense Network across different degradation models and datasets, emphasizing its effectiveness and robustness in image super-resolution tasks.

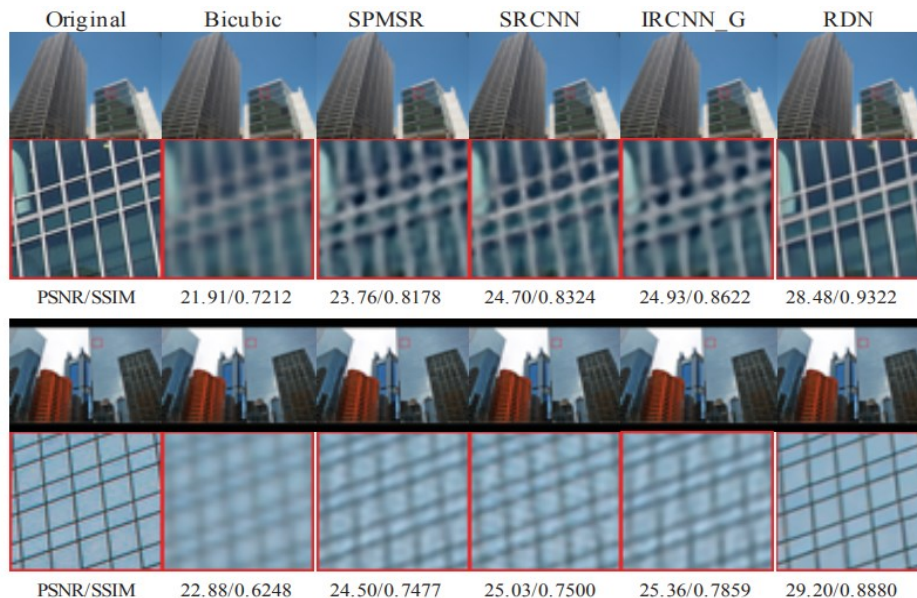


Fig. 12. Visual results using BD degradation model with scaling factor $\times 3$. The SR results are for image “img 096” from Urban100 and “img 099” from Urban100 respectively.

This method introduces a very deep Residual Dense Network (RDN) for image super-resolution (SR). The key building block of this network is the Residual Dense Block (RDB), where dense connections between layers enable the full utilization of local features. Local Feature Fusion (LFF) stabilizes training in a wider network and controls information preservation across RDBs. The RDB includes a Contiguous Memory (CM) mechanism, establishing direct connections between preceding RDBs and each layer in the current block. Local Residual Learning (LRL) enhances information and gradient flow. Additionally, the authors propose Global Feature Fusion (GFF) to extract hierarchical features in the low-resolution (LR) space. Through the integration of local and global features, the RDN achieves dense feature fusion and deep supervision. The proposed RDN is designed to handle three degradation models and real-world data. Extensive benchmark evaluations demonstrate the superiority of RDN over state-of-the-art methods in image super-resolution tasks.

Feedback Network for Image Super-Resolution [10]:

This method introduce a novel image super-resolution feedback network (SRFBN) with a feedback mechanism, utilizing top-down feedback flows through connections. The proposed network includes a feedback block (FB) that efficiently handles feedback information, enhances high-level representations, and incorporates up- and downsampling layers along with dense skip connections (Figure 13). Additionally, a curriculum-based training strategy is proposed, feeding high-resolution images with increasing reconstruction difficulty into the network in consecutive iterations. This approach allows the network to learn complex degradation models gradually, which is not possible with methods relying on only one-step prediction.

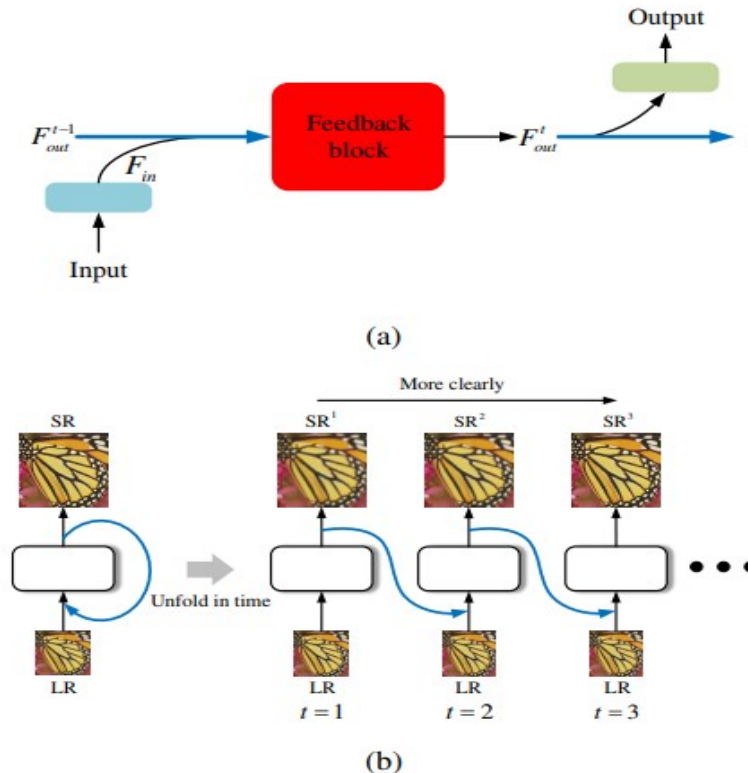


Fig. 13. The illustrations of the feedback mechanism in the proposed network. Blue arrows represent the feedback connections. (a) Feedback via the hidden state at one iteration. The feedback block (FB) receives the information of the input F_{in} and hidden state from last iteration F_{out}^{t-1} , and then passes its hidden state F_{out}^t to the next iteration and output. (b) The principle of our feedback scheme.

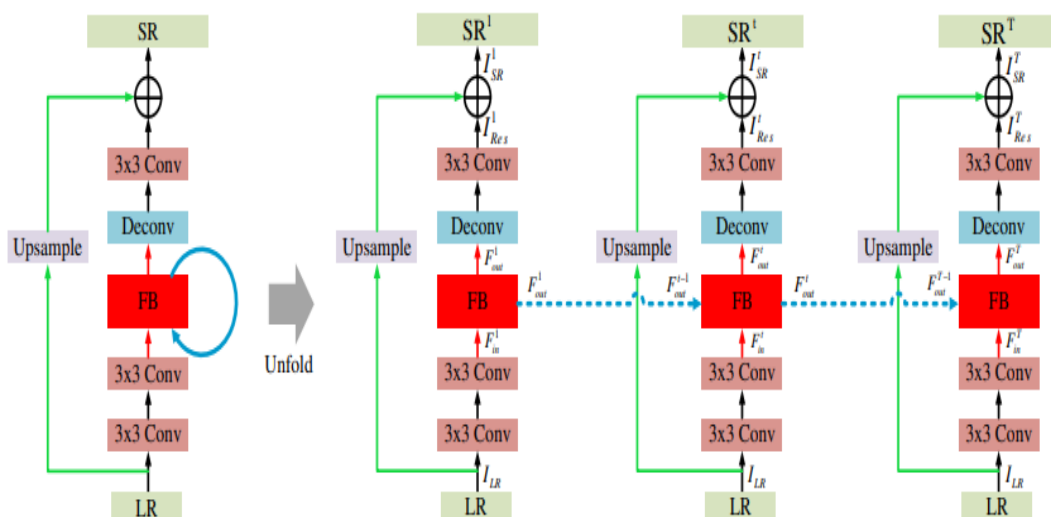


Fig. 14. . The architecture of our proposed super-resolution feedback network (SRFBN). Blue arrows represent feedback connections. Green arrows represent global residual skip connections.

The proposed network incorporates a feedback mechanism for correcting previous states based on output understanding, a feature less explored in image super-resolution (SR). Existing approaches in SR mainly follow a feedforward information flow. While related work used a convolutional recurrent neural network for feedback, it primarily addressed high-level vision tasks. In contrast, the authors introduce a feedback block (FB) in their super-resolution feedback network (SRFBN), demonstrating superior performance over ConvLSTM in information flow across hierarchical layers through dense skip connections. This makes SRFBN more fitting for image super-resolution tasks (Figure 14).

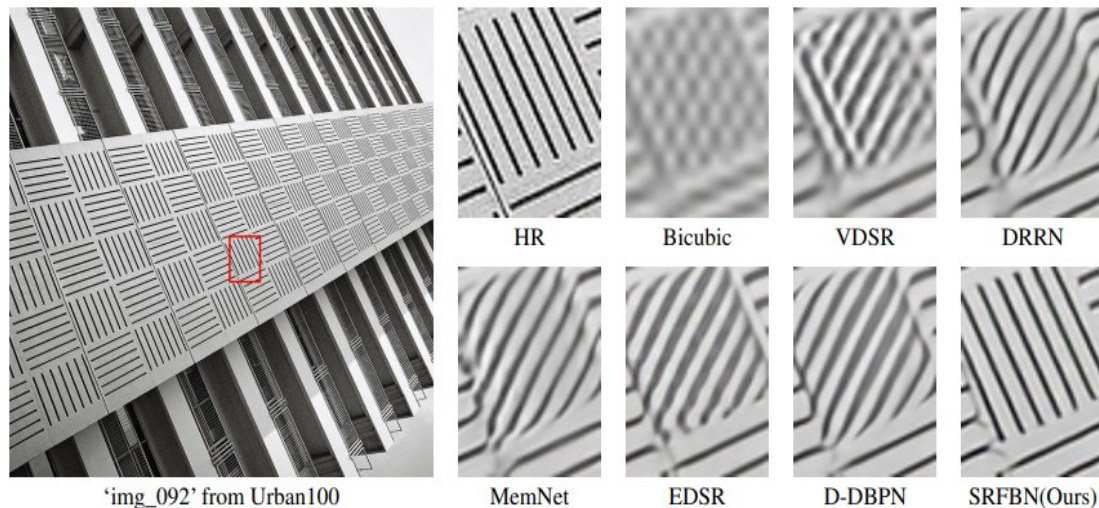


Fig. 15. Visual results of BI degradation model with scale factor $\times 4$.

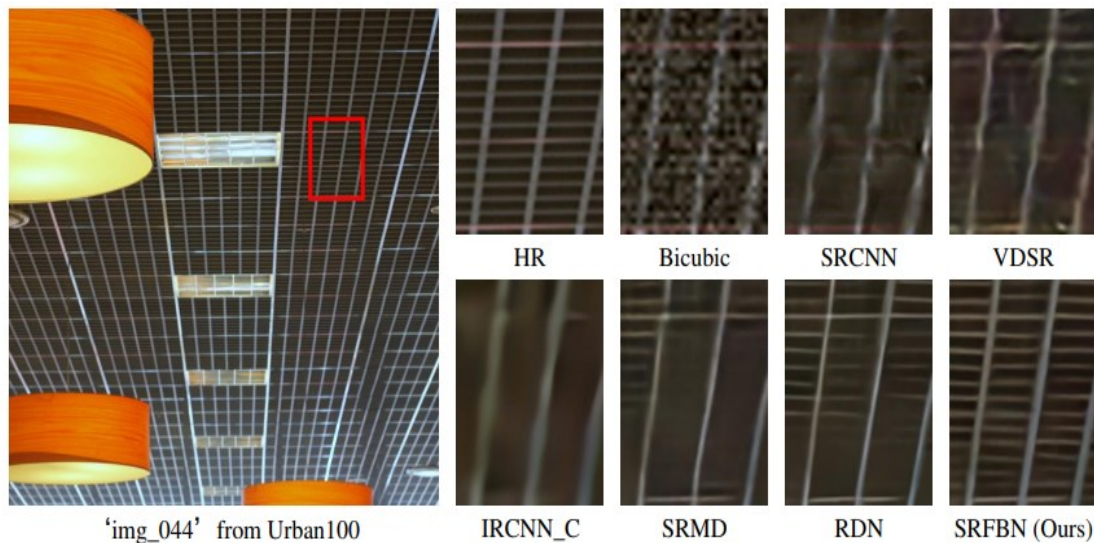


Fig. 16. Visual results of BD and DN degradation models with scale factor $\times 3$. The first set of images shows the results obtained from BD degradation model. The second set of images shows the results from DN degradation model.

BI Degradation Model: SRFBN outperforms state-of-the-art methods in image super-resolution, even against models with more resources like EDSR and D-DBPN, showcasing its superior performance (Figure 15).

BD and DN Degradation Models: SRFBN, trained with a curriculum learning strategy, consistently achieves the best quantitative results across BD and DN degradation models, demonstrating robustness and effectiveness in handling diverse degradation scenarios (Figure 16).

This method introduces SRFBN, a novel super-resolution feedback network that enhances low-level representations with high-level information to faithfully reconstruct SR images. The feedback block (FB) effectively manages feedback information and feature reuse. A curriculum learning strategy is proposed for handling complex degradation models in low-resolution images. Experimental results show that SRFBN achieves comparable or superior performance to state-of-the-art methods with significantly fewer parameters.

Single Image Super-Resolution via a Holistic Attention Network [11]:

In evaluating the Blur-downscale Degradation (BD) Model, the proposed method, Hierarchical Attention Network (HAN) and its

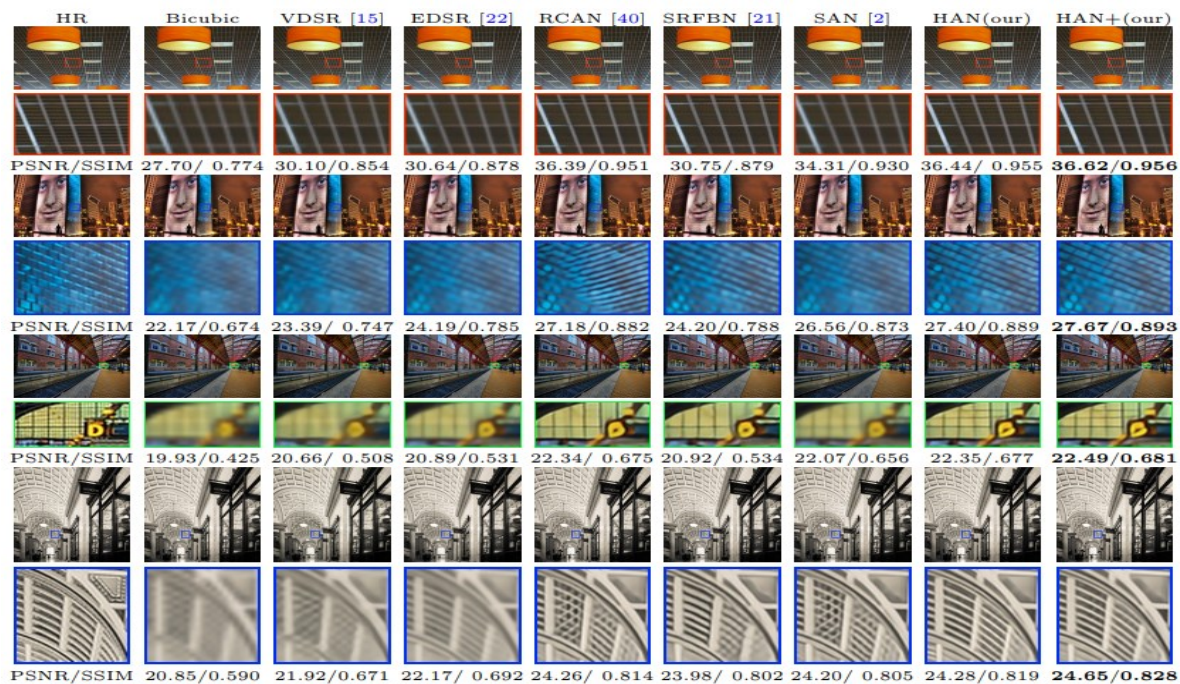


Fig. 17. Visual comparison for 3× SR with BD model on the Urban100 dataset. The best results are highlighted enhanced version HAN+, are compared with nine state-of-the-art super-resolution methods. Quantitative results for 3× super-resolution demonstrate that both HAN and HAN+ outperform existing methods, with HAN+ achieving the best scores. Visual results on Urban 100 dataset illustrate that HAN effectively recovers missing details in low-resolution images, surpassing other methods like VDSR, EDSR, RCAN, SRFBN, and SAN, which struggle with blurring artifacts and fail to recover structured details. HAN demonstrates its ability to suppress artifacts and exploit scene details for high-quality super-resolution. The method introduces a holistic attention network for single image super-resolution, utilizing the self-attention mechanism to adaptively learn global dependencies across depths, channels, and positions. The layer attention module captures long-distance dependencies among hierarchical layers, while the channel-spatial attention module incorporates channel and contextual information in each layer. These attention modules collaboratively enhance multi-level features, resulting in more informative features. Extensive experimental results on benchmark datasets showcase the proposed model's favorable performance against state-of-the-art super-resolution algorithms, excelling in accuracy and visual quality.

6. RESULTS & DISCUSSION

1. Comparative Evaluation of Super-Resolution Techniques

The review conducted a comparative analysis of prominent super-resolution techniques, including SRCNN, VDSR, EDSR, SRFBN, and HAN+. Quantitative metrics, such as PSNR and SSIM, were considered alongside visual assessments. Notably, HAN+ demonstrated superior performance, surpassing other methods in both quantitative measures and perceptual quality.

2. Evolution of Attention Mechanisms in Super-Resolution

The study highlighted the evolving role of attention mechanisms in super-resolution models. Specifically, models like HAN leveraged self-attention mechanisms to adaptively capture global dependencies across different depths, channels, and positions. The incorporation of attention mechanisms proved effective in enhancing the efficiency and overall performance of super-resolution techniques.

3. Addressing Challenges and Future Directions

Challenges identified in the review include real-time processing constraints, variability in image content, and model robustness to diverse degradation scenarios. The discussion underscores potential future directions, such as the integration of machine learning algorithms, more efficient attention mechanisms, and the exploration of domain-specific super-resolution solutions.

4. Practical Implications and Applications

The findings emphasize the practical implications of super-resolution techniques in diverse applications, ranging from medical imaging to surveillance. The ability to enhance image quality has significant implications for data analysis and interpretation in various real-world scenarios.

7. CONCLUSION

This review provides a comprehensive examination of various super-resolution techniques employed to enhance image resolution. Through a comparative analysis of state-of-the-art methods, it was observed that models such as HAN+ exhibit superior performance, emphasizing the importance of evolving techniques in achieving higher accuracy and perceptual quality.

The integration of attention mechanisms, notably self-attention in models like HAN, reflects a growing trend in adapting super-resolution techniques to better capture global dependencies across image depths, channels, and positions. This not only enhances the overall efficiency of these methods but also opens avenues for further advancements in the field.

Despite the progress made, challenges persist, including real-time processing limitations and the need for enhanced robustness to diverse degradation scenarios. The identified challenges, along with potential future directions discussed in the review, offer valuable insights for researchers seeking to advance the field of image resolution through super-resolution.

The practical implications of super-resolution techniques extend to various domains, with the potential to revolutionize applications such as medical imaging and surveillance. The ability to improve image quality has far-reaching implications for data analysis and interpretation, underscoring the practical significance of advancements in this field.

In essence, this review consolidates current knowledge in the realm of image resolution through super-resolution techniques, providing a foundation for future research directions. The ongoing evolution of methodologies, attention mechanisms, and practical applications signifies the dynamic nature of this field, urging continued exploration and innovation in the pursuit of enhanced image resolution.

8. REFERENCES

- [1] Chitwan Saharia, Jonathan Ho, William Chan , Tim Salimans, David J. Fleet , and Mohammad Norouzi. (2023). Image Super-Resolution via Iterative Refinement.
- [2] Juncheng Li1, Faming Fang1, Kangfu Mei, and Guixu Zhang. (2018). Multi-scale Residual Network for Image Super-Resolution.
- [3] Tao Dai1, Jianrui Cai3, Yongbing Zhang1, Shu-Tao Xia1, Lei Zhang. Second-order Attention Network for Single Image Super-Resolution.
- [4] Wei-Sheng Lai, Jia-Bin Huang, Narendra Ahuja, and Ming-Hsuan Yang. (2018). Fast and Accurate Image Super-Resolution with Deep Laplacian Pyramid Networks.
- [5] Yuan Yuan, Siyuan Liu, Jiawei Zhang, Yongbing Zhang, Chao Dong, Liang Lin. Unsupervised Image Super-Resolution using Cycle-in-Cycle Generative Adversarial Networks.
- [6] Yulun Zhang, Kunpeng Li, Kai Li, Lichen Wang, Bineng Zhong, Yun Fu. (2018). Image Super-Resolution Using Very Deep Residual Channel Attention Networks.
- [7] Zheng Hui, Xiumei Wang, Xinbo Gao. Fast and Accurate Single Image Super-Resolution via Information Distillation Network.
- [8] Zheng Hui, Xinbo Gao, Yunchu Yang, Xiumei Wang. (2019). Lightweight Image Super-Resolution with Information Multi-distillation Network.
- [9] Yulun Zhang, Yapeng Tian, Yu Kong, Bineng Zhong, Yun Fu. Residual Dense Network for Image Super-Resolution.
- [10] Zhen Li, Jinglei Yang, Zheng Liu, Xiaomin Yang, Gwanggil Jeon, Wei Wu. Feedback Network for Image Super-Resolution.
- [11] Ben Niu, Weilei Wen, Wenqi Ren, Xiangde Zhang, Lianping Yang, Shuzhen Wang, Kaihao Zhang, Xiaochun Cao, Haifeng Shen. (2020). Single Image Super-Resolution via a Holistic Attention Network.
- [12] Saeed Anwar, Salman Khan, and Nick Barnes. (2020). A Deep Journey into Super-resolution: A Survey.