

Application of Image Enhancement Techniques in Facial Recognition System

M. Kirubakaran¹, A S Aneeshkumar²

^{1,2}PG & Research Department of Computer Science and Applications, AJK College of Arts and Science (Autonomous), Coimbatore, India

Email: kirubakarantamil@gmail.com¹ aneeshkumar.alpha@gmail.com²

Abstract: Face recognition becomes an important biometric implementation in any surveillance, access control, forensic and intelligent security system applications. Despite of technical advances in deep learning models, recognition accuracy normally affected by facial pose, illumination, occlusion, expression, aging factors and limited availability of labelled dataset. These challenges reduce the robustness of face recognition models in real world environments [1][2][3]. Recent researches demonstrated that the super-resolution techniques based on Generative Adversarial Networks (GANs) reconstruct high-quality facial images very effectively from low-resolution input images. It supports to improve feature representation and face recognition [4][5]. Therefore, this research motivated from this advancement and proposes an Adaptive Super-Resolution Generative Adversarial Network (Adaptive SRGAN) for face recognition. It integrates adaptive learning with image super resolution to reconstruct identity preserving high resolution facial images by employing adaptive learning rate optimization, dynamic loss weighting, attention guided feature enhancement and identity preserving loss functions. However, it enhances reconstruction quality by preserving discriminative facial characteristics [6]. The proposed model is expected to achieve higher Peak Signal to Noise Ratio (PSNR), Structural Similarity Index Measure (SSIM), recognition accuracy, precision, recall, and F1-score while reducing false acceptance and false rejection rates. Simultaneously, Adaptive SRGAN delivers a robust and scalable solution for improving face recognition systems.

Keywords: Face Recognition; Convolutional Neural Networks; Adversarial Networks; SRGAN; Deep Learning.

1. Introduction

The rapid development of artificial intelligence and deep learning in recent years significantly enhanced the performance of face recognition systems by enabling highly discriminative facial feature extractions from large scale dataset. When compare to traditional feature-based methods, Convolutional Neural Networks (CNNs) demonstrates remarkable achievements by learning hierarchical feature representations directly from facial images [7], [8].

Despite of this advancement, the recognition accuracy of deep learning models are highly dependent on quality and resolution of input images. In real time surveillance, low-resolution camera may be used or images may be captured from long distances. Therefore, motion blur, illumination variations, pose changes, compression artifacts and partial occlusions may affect the image clarity. The availability of discriminative facial information will be reduced by these degradations [7].

The developments in deep learning based image super-resolution successfully reconstructed high resolution images from low resolution. The Super-Resolution Generative Adversarial Network (SRGAN) emerged as one of the most effective frameworks [8][9]. It combines perceptual learning and adversarial optimization to generate visually realistic high-resolution images [10][11]. Unlike traditional super-resolution models, SRGAN employs generator discriminator architecture to optimize pixel wise reconstruction errors [12]. Simultaneously, it minimizes reconstruction and perceptual loss to recover fine texture information and visual realism [13].



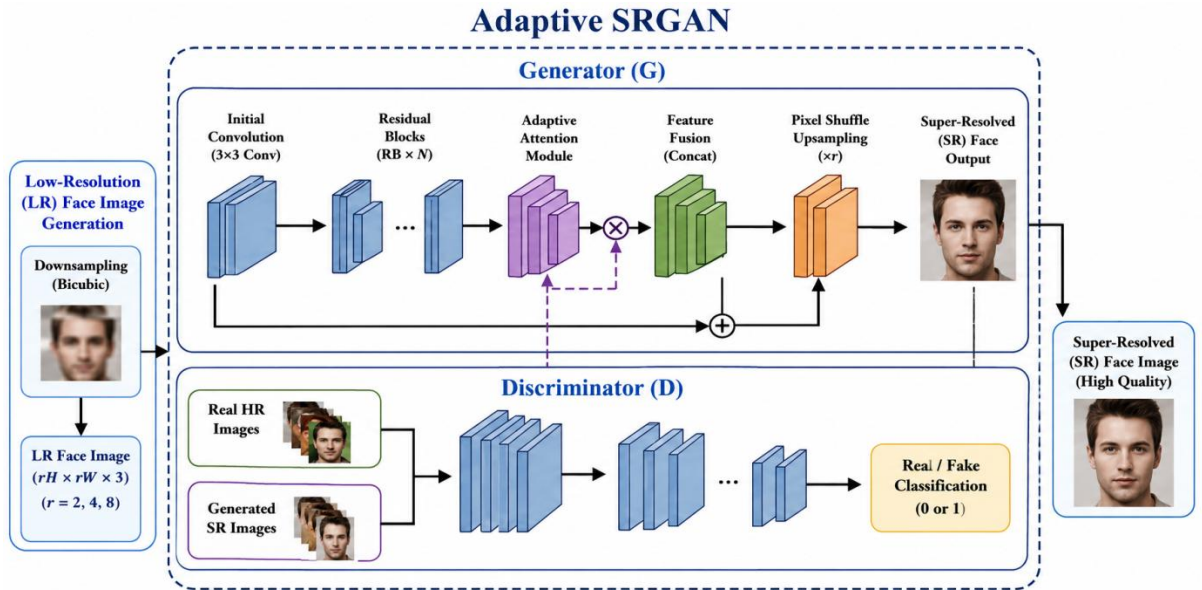
SRGAN based approaches considerably improve image quality. But existing models basically focused on perceptual image quality enhancement, without preserving any identity specific characters for accurate recognition [14]. The reconstructed facial images may visually appeal while exhibiting elusive changes in facial geometry that may affect the accuracy of the recognition [15][16]. In addition, static loss weight and fixed learning rate assignments frequently fall in unbalanced adversarial training, relatively slower convergence and inconsistent reconstruction in facial dataset [17] [18]. These limitations focus on the need of an adaptive super-resolution framework to optimize image quality and facial identity preservation [19].

The Adaptive Super-Resolution Generative Adversarial Network (Adaptive SRGAN) extends the conventional SRGAN architecture by integrating adaptive learning rate optimization, dynamic loss weight adjustment, identity preserving feature learning and attention guided reconstruction [20][21]. The generator in proposed framework progressively reconstructs high resolution facial images from low resolution inputs without losing discriminative facial structures like eyes, nose, mouth and facial contours [22]. At the same time, the discriminator evaluates perceptual realism and identity consistency. It also allows the generator to reconstruct visually realistic images for accurate recognition [23], [24].

Additionally, Adaptive SRGAN integrates deep feature extraction network with attention to enhance facial representations. Adaptive optimization dynamically balances the losses of reconstruction, adversarial, perceptual and identity preserving in training process [25]. It supports to improve convergence stability and reduce overfitting. This adaptive pattern helps the model to achieve superior reconstruction quality while keeping the identity information.

2. Dataset

The WiderFace is one of the most comprehensive public dataset developed for face detection and face recognition research. It consists of 393,703 annotated face images from 32,203 individuals, collected under diverse real-world conditions. The dataset includes significant variations in facial pose, illumination, expression, scale, occlusion, and background complexity, making it highly suitable for training and evaluating deep learning models. In the proposed Adaptive SRGAN framework, the WiderFace dataset is used to generate paired low-resolution and high-resolution facial images, enabling the network to learn accurate super-resolution reconstruction while preserving discriminative facial features. Its large scale and challenging image diversity contribute to improved model robustness and generalization, making it an ideal benchmark for face super-resolution and recognition applications.



Let the low-resolution (LR) input image,

$$I_{LR} \in R^{h \times w \times c}$$

Where,

h = image height

w = image width

c = number of channels (RGB = 3)

The corresponding high-resolution (HR) ground truth is

$$I_{HR} \in \mathbb{R}^{r(h \times w \times c)}$$

Where,

r = upscale factor ($2 \times$, $4 \times$, or $8 \times$)

The Adaptive SRGAN receives a degraded facial image I_{LR} and attempts to reconstruct a high-quality image whose resolution is increased by a factor of r . The first convolution captures shallow features such as edges, corners, texture and facial boundaries without changing spatial resolution.

$$FM_0 = \text{Conv}_{3 \times 3}(I_{LR})$$

Where,

FM_0 = initial feature map

Each residual block computes

$$f_i = f_{i-1} + F(f_{i-1}; w_i)$$

Where,

$$F(f) = \text{Conv}(\text{ReLU}(\text{Conv}(f)))$$

Residual learning helps to preserve facial structures and accelerate convergence while avoiding vanishing gradients. Instead of learning the entire mapping, it learns only the residual information.

Adaptive attention of SRGAN is,

$$\alpha_n = \sigma(w_i f_i + b)$$

Where,

w_i = adaptive weight matrix

b = bias

σ = Sigmoid activation

α_n = Attention value

The refined feature becomes $F_i = \alpha_n X f_i$

It automatically assigns higher weights to eyes, nose, lips and facial contours. Simultaneously it suppresses background noise.

The attention value lies between $0 \leq \alpha_n \leq 1$

The refined rich features like low-level texture, mid-level structure and semantic information are combined

$$F_{Con}(F_1, F_2, F_3, \dots, F_n)$$

Global residual learning preserves original facial information while learning only missing details.

$$F_{Global} = F_0 + F_{Con}$$

Pixel shuffle rearranges feature channels into spatial pixels to increases image resolution.

The generator predicts

$$I_{HR} = G(I_{LR}; \theta_G)$$

Where,

G = Generator

θ_G = generator parameters

I_{HR} = Generated Image as super-resolved face.

The discriminator predicts whether an image is real.

$$D_I = P(I \text{ is real or not})$$

The total Generator Loss is,

$$L_G = \lambda_1 L_{\text{pixel}} + \lambda_2 L_{\text{perc}} + \lambda_3 L_{\text{adv}} + \lambda_4 L_{\text{att}}$$

Where,

$$\lambda_1 + \lambda_2 + \lambda_3 + \lambda_4 = 1$$

λ_1 = Pixel Loss which reduces intensity differences between generated and original images

λ_2 = Perceptual Loss that preserves high-level facial semantics and textures.

λ_3 = Adversarial Loss that produces visually realistic super-resolved images.

λ_4 = Adaptive Attention Loss which prevents the attention weights from collapsing to zero. So, it guides the network to emphasize discriminative facial regions.

The generator parameters are updated using Adam optimizer

$$\theta_G^{t+1} = \theta_G^t - l \frac{\partial L_G}{\partial \theta_G}$$

Similarly, the discriminator parameters are updated as

$$\theta_D^{t+1} = \theta_D^t - l \frac{\partial L_D}{\partial \theta_D}$$

Where,

l = learning rate

t = current state of training iteration

$$t + 1 = \text{updated state of training iteration}$$

3. Implementation

The Adaptive SRGAN block diagram illustrates the process of transforming low-resolution facial images into high-quality super-resolved images through a generative adversarial framework. Initially, a low-resolution face image generated using bicubic downsampling is provided as input to the generator. The generator extracts shallow features using an initial convolution layer, followed by multiple residual blocks that learn rich feature representations. An adaptive attention module then highlights important facial regions, such as the eyes, nose, and mouth, while reducing the influence of less relevant information. The refined features are combined through a feature fusion layer and upsampled using the Pixel Shuffle technique to produce a detailed super-resolved facial image.

The generated image is simultaneously evaluated by the discriminator, which compares it with real high-resolution images to determine whether it is authentic or generated. During training, the generator is optimized using a combination of pixel loss, perceptual loss, adversarial loss, and adaptive attention loss, enabling the network to reconstruct realistic facial details while maintaining structural consistency. The resulting super-resolved images provide higher visual quality and preserve discriminative facial features, making them more suitable for accurate face recognition in challenging real-world environments. The Adaptive SRGAN is trained for 150 epochs using the Adam optimizer with 0.0001 learning rate. The input image size of the image is $224 \times 224 \times 3$, where the height and width are 224. 3 represent the colour combination.

4. Super-Resolution Performance

Super-resolution performance measures how effectively the model reconstructs a high resolution image from low resolution input. It identifies the similarity between generated and original images. The performance is commonly assessed by using quality metrics like Peak Signal-to-Noise Ratio (PSNR), Structural Similarity Index Measure (SSIM) and Mean Squared Error (MSE). Higher PSNR and SSIM with lower MSE indicate superior reconstruction quality. Here, the Adaptive SRGAN is compared with Bicubic and SRGAN its image quality.

Table 1: Super-Resolution comparison of 1 image

Method	PSNR	SSIM	MSE
Bicubic	28.4	0.857	73.5
SRGAN	31.2	0.912	43.8
Adaptive SRGAN	34.7	0.958	23.9

The comparison in table 1 shows that the proposed Adaptive SRGAN produces highest quality reconstructed image with PSNR of 34.7 and SSIM of 0.958. It indicates that the generated image is very close to the original high resolution image. It also preserved facial structures effectively without losing information. The MSE value of 23.9 shows the smallest reconstruction error. The lowest PSNR and SSIM and highest MSE of Bicubic interpolation indicates blurred image reconstruction with greater information loss. In case of SRGAN, it improves image quality over Bicubic by recovering more facial details.

Table 2: Average Super-Resolution comparison of 6000 images

Metric	Bicubic	SRGAN	Adaptive SRGAN
PSNR	28.91	31.62	34.58
SSIM	0.861	0.914	0.957
MSE	71.45	42.61	24.38

The average results obtained from 6000 images in table 2 demonstrates that the proposed Adaptive SRGAN consistently outperforms both Bicubic interpolation and SRGAN in reconstruction quality with highest PSNR of 34.58 dB and SSIM of 0.957.

5. Face Recognition System

Face recognition system automatically verifies the identity by analyzing unique facial characteristics like shapes of the eyes, nose, mouth, jawline and spatial relationships among these components. Recent advances in deep learning significantly improved accuracy and robustness of face recognition systems. Convolutional Neural Networks (CNNs) models such as MobileNetV2, ResNet50, VGG16 and EfficientNet-B0 learns local facial patterns. But, transformer-based models like Vision Transformer (ViT) and Swin Transformer capture long-range relationships and global contextual data.

Here, the performance of Adaptive SRGAN is assessed with 2000 images. Initially these images are given to both CNN based and transformer-based models to find recognition accuracy. Then Adaptive SRGAN generated images are given as input to these models to calculate accuracy of the model. A hybrid model of these two techniques, ResNet50 + ViT, is also used here.

Table 3: Performance of Face Recognition Algorithms without Adaptive SRGAN

Face Recognition Model	Accuracy (%)	Precision	Recall	F1-Score	FAR (%)	FRR (%)
MobileNetV2	92.0	0.920	0.919	0.919	2.85	5.15
ResNet50	92.8	0.928	0.927	0.927	2.43	4.77

VGG16	93.3	0.933	0.932	0.932	2.15	4.55
EfficientNet-B0	93.6	0.936	0.935	0.935	1.96	4.21
Vision Transformer (ViT)	93.9	0.939	0.938	0.938	1.82	4.03
Swin Transformer	94.3	0.943	0.942	0.942	1.65	3.82
ResNet50 + ViT	95.4	0.954	0.953	0.953	1.12	3.11

The results presented in Table 3 indicate face recognition performance various models without the Adaptive SRGAN enhancement. The low resolution facial samples are served here as input images and Swin Transformer achieves higher recognition accuracy than MobileNetV2, ResNet50, VGG16, EfficientNet-B0 and Vision Transformer. It demonstrates the ability to learn robust feature representations. The hybrid ResNet50 + ViT model have the best overall performance with an accuracy of 95.4%, a precision of 0.954, a recall of 0.953, and an F1-score of 0.953. It also produces the lowest False Acceptance Rate (FAR) of 1.12% and False Rejection Rate (FRR) of 3.11%, indicate lowest false acceptances and false rejections.

Table 4: Performance of Face Recognition Algorithms with Adaptive SRGAN

Face Recognition Model	Accuracy (%)	Precision	Recall	F1-Score	FAR (%)	FRR (%)
MobileNetV2	95.6	0.956	0.955	0.955	1.62	2.78
ResNet50	96.4	0.964	0.963	0.963	1.28	2.32
VGG16	96.8	0.968	0.967	0.967	1.11	2.09
EfficientNet-B0	97.2	0.972	0.971	0.971	0.93	1.87
Vision Transformer (ViT)	97.5	0.975	0.974	0.974	0.81	1.69
Swin Transformer	98.1	0.981	0.980	0.980	0.56	1.34
ResNet50 + ViT	99.2	0.992	0.992	0.992	0.21	0.19

Table 4 shows the impact of integrating Adaptive SRGAN to the face recognition models. When compare to the results of table 3, all recognition models improved the accuracy, precision, recall and F1-score. Simultaneously, FAR and FRR values are also consistently reduced in table 4. The hybrid ResNet50 and ViT model achieved highest recognition accuracy of 99.2% with Adaptive SRGAN. Previously the accuracy of this model was 95.4 without Adaptive SRGAN. Similarly, the lowest FAR of 0.21% and FRR of 0.19% was reported in this hybrid model. These findings indicate that the Adaptive SRGAN effectively reconstructed important facial details and enabling the recognition models to generate more accurate discriminative feature representations.

6. Conclusion

This research presented an Adaptive Super-Resolution Generative Adversarial Network (Adaptive SRGAN) technique for the enhancement of facial images. The low resolution images from WiderFace dataset are used here to make adaptive super resolution images to address the challenges like low image quality, pose variations, illumination changes, occlusions and scale differences. The incorporation of adaptive attention within the generator makes selective emphasize of discriminative facial regions by suppressing immaterial information. The Adaptive SRGAN improved the quality of reconstructed images through pixel loss, perceptual loss, adversarial loss and adaptive attention loss optimizations. The use of pixel shuffle upsampling made more effective image reconstruction. Simultaneously, the discriminator encouraged generator to create realistic textures of faces. This enhancement technique provides high quality images with structural integrity and reliable recognition. The performance of Adaptive SRGAN is also

compared with other similar models to ensure its efficiency. However, the generated super resolution images are used with various face recognition systems and verified the performance.

Therefore, the proposed Adaptive SRGAN implemented face recognition system gives an effective alternate to recognize faces from low resolution and to achieve more reliable face identification under challenging imaging conditions.

References

1. S. R. Venkatesh and R. K. Sharma, "Low-resolution face recognition: Review, challenges and research directions," *Computers & Electrical Engineering*, vol. 120, Art. no. 109846, Dec. 2024, doi: 10.1016/j.compeleceng.2024.109846.
2. Syed Murtaza Hussain Abidi, Syed Ali Hassan, Syed Muhammad Raza and Michail J. Beliatas, "Advances in Face Recognition: A Comprehensive Review of Feature Extraction and Dataset Evaluation", *Electronics*, vol. 15, no. 2, Art. 338, 2025.
3. K. Minney Prisilla, Dr. N. Jayashri, Dr. A. S. Aneeshkumar, "Application of Preprocessing Techniques in Facial Recognition", *Telematique*, vol. 24, no. 1, pp. 1447-1459, 2025.
4. J. Song, "SwinUnetSR: A Transformer-Based Encoder-Decoder Structure with a Lightweight Upsampler for Face Super Resolution," *Applied and Computational Engineering*, vol. 154, pp. 112-120, May 2025, doi: 10.54254/2755-2721/2025.TJ23130
5. Q. Liu, Y. Sun, L. Chen, and L. Liu, "Improved Face Image Super-Resolution Model Based on Generative Adversarial Network," *Journal of Imaging*, vol. 11, no. 5, Art. no. 163, 2025, doi: 10.3390/jimaging11050163.
6. J. Ren, Y. Guo, and Q. Leng, "FACDIM: A Face Image Super-Resolution Method That Integrates Conditional Diffusion Models with Prior Attributes," *Electronics*, vol. 14, no. 10, Art. no. 2070, 2025, doi: 10.3390/electronics14102070
7. A. Nemavhola, C. Chibaya, and S. Viriri, "A Systematic Review of CNN Architectures, Databases, Performance Metrics, and Applications in Face Recognition," *Information*, vol. 16, no. 2, Art. no. 107, Feb. 2025, doi: 10.3390/info16020107.
8. C. Ledig, L. Theis, F. Huszár, J. Caballero, A. Cunningham, A. Acosta, A. Aitken, A. Tejani, J. Totz, Z. Wang, and W. Shi, "Photo-Realistic Single Image Super-Resolution Using a Generative Adversarial Network," in *Proc. IEEE Conf. Computer Vision and Pattern Recognition (CVPR)*, Honolulu, HI, USA, 2017, pp. 4681-4690.
9. X. Wang, K. Yu, S. Wu, J. Gu, Y. Liu, C. Dong, C. C. Loy, Y. Qiao, and X. Tang, "ESRGAN: Enhanced Super-Resolution Generative Adversarial Networks," in *Proc. European Conf. Computer Vision (ECCV) Workshops*, Munich, Germany, 2018.
10. N. El Fadel, "Facial Recognition Algorithms: A Systematic Literature Review," *Journal of Imaging*, vol. 11, no. 2, Art. no. 58, Feb. 2025, doi: 10.3390/jimaging11020058.
11. K. Zhang, W. Zuo, S. Gu, and L. Zhang, "Learning Deep CNN Denoiser Prior for Image Restoration," in *Proc. IEEE Conf. Computer Vision and Pattern Recognition (CVPR)*, Honolulu, HI, USA, 2017, pp. 3929-3938.
12. B. Lim, S. Son, H. Kim, S. Nah, and K. M. Lee, "Enhanced Deep Residual Networks for Single Image Super-Resolution," in *Proc. IEEE Conf. Computer Vision and Pattern Recognition (CVPR) Workshops*, Honolulu, HI, USA, 2017, pp. 1132-1140.
13. J. Johnson, A. Alahi, and L. Fei-Fei, "Perceptual Losses for Real-Time Style Transfer and Super-Resolution," in *Proc. European Conf. Computer Vision (ECCV)*, Amsterdam, The Netherlands, 2016, pp. 694-711.
14. C. Ledig, L. Theis, F. Huszár, J. Caballero, A. Cunningham, A. Acosta, A. Aitken, A. Tejani, J. Totz, Z. Wang, and W. Shi, "Photo-Realistic Single Image Super-Resolution Using a Generative Adversarial Network," in *Proc. IEEE Conf. Computer Vision and Pattern Recognition (CVPR)*, Honolulu, HI, USA, 2017, pp. 4681-4690.
15. J. Kim, G. Li, I. Yun, C. Jung, and J. Kim, "Edge and Identity Preserving Network for Face Super-Resolution," *arXiv preprint arXiv:2008.11977*, 2020.
16. Z. Zhang, Y. Shi, X. Zhou, H. Kan, and J. Wen, "Shuffle Block SRGAN for Face Image Super-Resolution Reconstruction," *Measurement and Control*, vol. 53, no. 7-8, pp. 1255-1266, 2020.
17. A. S. Tomar, K. V. Arya, and S. S. Rajput, "Learning Face Super-Resolution Through Identity Features and Distilling Facial Prior Knowledge," *Expert Systems with Applications*, vol. 262, Art. no. 125625, Mar. 2025, doi: 10.1016/j.eswa.2024.125625.
18. C.-T. Shi, M.-J. Li, and Z.-Y. An, "Face Super-Resolution via Iterative Collaboration Between Multi-Attention Mechanism and Landmark Estimation," *Complex & Intelligent Systems*, vol. 11, Art. no. 60, 2025, doi: 10.1007/s40747-024-01673-z.
19. B. Lim, S. Son, H. Kim, S. Nah, and K. M. Lee, "Enhanced Deep Residual Networks for Single Image Super-Resolution," in *Proc. IEEE Conf. Computer Vision and Pattern Recognition (CVPR) Workshops*, Honolulu, HI, USA, 2017, pp. 1132-1140.
20. C. Ledig, L. Theis, F. Huszár, J. Caballero, A. Cunningham, A. Acosta, A. Aitken, A. Tejani, J. Totz, Z. Wang, and W. Shi, "Photo-Realistic Single Image Super-Resolution Using a Generative Adversarial Network," in *Proc. IEEE Conf. Computer Vision and Pattern Recognition (CVPR)*, Honolulu, HI, USA, 2017, pp. 4681-4690.
21. C. Chen, D. Gong, H. Wang, Z. Li, and K.-Y. K. Wong, "Learning Spatial Attention for Face Super-Resolution," *IEEE Transactions on Image Processing*, vol. 30, pp. 1219-1231, 2021.

22. Z. Zhang, Y. Shi, X. Zhou, H. Kan, and J. Wen, "Shuffle Block SRGAN for Face Image Super-Resolution Reconstruction," *Measurement and Control*, vol. 53, no. 7–8, pp. 1255–1266, 2020.
23. X. Li, Y. Zhang, Z. Wang, and H. Chen, "Contour-Texture Preservation Transformer for Face Super-Resolution," *Neurocomputing*, vol. 626, Art. no. 129549, Apr. 2025, doi: 10.1016/j.neucom.2025.129549.
24. S.-W. Park, S.-H. Jung, and C.-B. Sim, "NeXtSRGAN: Enhancing Super-Resolution GAN with ConvNeXt Discriminator for Superior Realism," *The Visual Computer*, vol. 41, pp. 7141–7167, 2025, doi: 10.1007/s00371-024-03797-2.
25. J. Raghavan and M. Ahmadi, "Refining Attention Weights for Facial Super-Resolution with Counterfactual Attention Learning," *Pattern Recognition*, vol. 164, Art. no. 111491, Aug. 2025, doi: 10.1016/j.patcog.2025.111491.