



Elevating Image Resolution Using Deep And Shallow Network

T.SANDEEP, K.HIMA GANESH, T.AYYAPPA, R.BALA VENKATA SAI, P.SIVAKRISHNA

STUDENT, STUDENT, STUDENT, STUDENT, ASSISTANT PROFESSOR
DEPARTMENT OF INFORMATION TECHNOLOGY,
SESHADRI RAO GUDLAVALLERU ENGINEERING
COLLEGE, GUDLAVALLERU, ANDHRA PRADESH, INDIA

Abstract

Image Super-Resolution (ISR) is a critical problem in computer vision aimed at reconstructing high-resolution (HR) images from degraded low-resolution (LR) counterparts. In this paper, we present a deep learning-based approach leveraging a Convolutional Neural Network (CNN) architecture known as Super-Resolution Convolutional Neural Network (SRCNN) to enhance image quality. The SRCNN model, composed of three convolutional layers, is trained on a large-scale dataset with pre-trained weights to improve the perceptual quality of LR images. The model achieves image enhancement by performing feature extraction, non-linear mapping, and reconstruction, primarily focusing on the Y (Luminance) channel of YCrCb color space to retain color fidelity. We evaluate the performance of the SRCNN model using standard metrics such as Peak Signal-to-Noise Ratio (PSNR), Structural Similarity Index (SSIM), and Mean Squared Error (MSE). Additionally, we compare the performance of images before and after enhancement using bar graphs for a comprehensive analysis.

Keywords: Image Super-Resolution, Convolutional Neural Network, SRCNN, PSNR, SSIM, Deep Learning.

Introduction

In recent years, the rapid growth of multimedia applications and advancements in imaging devices have heightened the demand for high-resolution (HR) images. High-quality images are essential across various domains, including medical imaging, remote sensing, surveillance, entertainment, and autonomous driving. However, capturing high-resolution images directly can be costly, time-consuming, or even impractical under certain conditions. Therefore, developing efficient algorithms for image super-resolution (ISR) has become a critical area of research. Image Super-Resolution is a technique aimed at reconstructing HR images from their degraded low-resolution (LR) counterparts. Traditional methods for super-resolution, such as interpolation techniques (bilinear, bicubic, and nearest-neighbor), often produce images with blurred edges and lack fine details. With the advent of deep learning, convolutional neural networks (CNNs) have shown remarkable success in numerous computer vision tasks, including image classification, object detection, segmentation, and super-resolution.

Among various deep learning architectures, the Super-Resolution Convolutional Neural Network (SRCNN) proposed by Dong et al. has proven to be an effective and efficient approach for ISR. The SRCNN model operates directly on the low-resolution image patches, learning an end-to-end mapping between the LR and HR images. Unlike traditional methods, which rely on hand-crafted features, SRCNN learns hierarchical features from raw pixels, making it highly adaptable to various datasets and degradation models.

The core of the SRCNN architecture involves three convolutional layers responsible for feature extraction, non-linear mapping, and reconstruction. The network is trained using a Mean Squared Error (MSE) loss function, optimized through the Adam optimizer. Additionally, to achieve a high-quality super-resolved image, the model is trained on the luminance channel (Y channel) of the YCrCb color space since it retains most of the structural information while maintaining color consistency.

This study aims to implement and evaluate an SRCNN model using pre-trained weights to enhance the quality of degraded images. Performance is assessed using standard metrics such as Peak Signal-to-Noise Ratio (PSNR), Structural Similarity Index (SSIM), and Mean Squared Error (MSE). To provide a comprehensive comparison, we present the results before and after super-resolution enhancement using bar graphs.

Methodology

The proposed methodology involves implementing an SRCNN model to enhance low-resolution images by learning a non-linear mapping between low-resolution and high-resolution image patches. This section details the model architecture, dataset preparation, training process, and evaluation metrics.

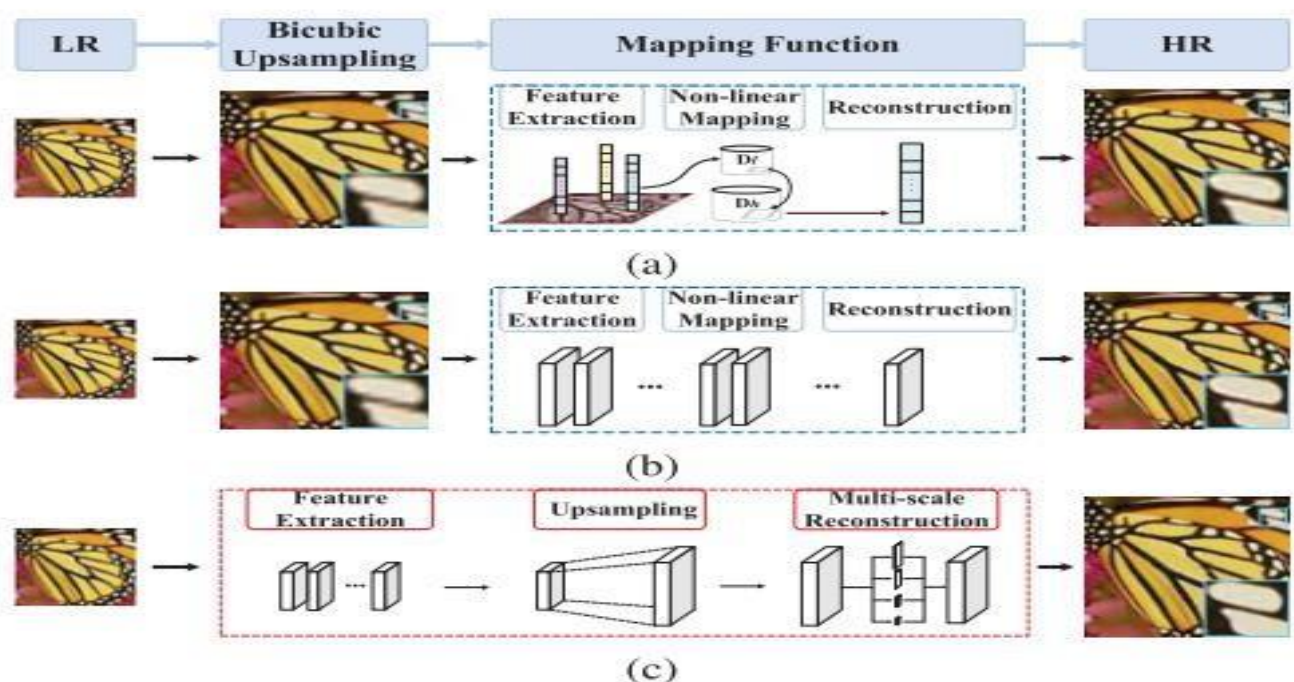
A. Model Architecture

The Super-Resolution Convolutional Neural Network (SRCNN) consists of three convolutional layers, each performing a specific task:

Feature Extraction: The first layer extracts features from the input image using 128 filters of size 9x9, applying ReLU activation. This layer captures low-level features and edges, producing feature maps.

Non-Linear Mapping: The second layer applies 64 filters of size 3x3 with ReLU activation. This layer maps the extracted features to high-resolution space using a non-linear transformation.

Reconstruction: The final layer uses a single filter of size 5x5 with linear activation to reconstruct the high-resolution image from the transformed feature maps.



B. Dataset Preparation

Images are degraded using bicubic downscaling and then upsampled to their original size to generate low-resolution samples. The model is trained and evaluated using a dataset of natural images, where each input image is degraded by a factor of 2 or 4.

C. Training Process

The model is trained using pre-trained weights (3051crop_weight_200.h5) to reduce training time and enhance performance. The training process involves:

- Feeding low-resolution images as input to the model.
- Calculating the loss between the generated high-resolution image and the original image.
- Updating the model weights using backpropagation.

D. Evaluation Metrics

The performance of the SRCNN model is evaluated using the following metrics:

1. Peak Signal-to-Noise Ratio (PSNR)

PSNR measures the **ratio between the maximum possible power of a signal and the power of the noise** that affects the fidelity of its representation. Higher PSNR indicates better image quality.

$$\text{PSNR} = 20 \cdot \log_{10} \left(\frac{\text{MAX}_I}{\sqrt{\text{MSE}}} \right)$$

Where:

MAX_I is the maximum possible pixel value of the image (e.g., 255 for 8-bit images).

MSE is the Mean Squared Error between the original image and the reconstructed image.

2. Mean Squared Error (MSE)

MSE measures the average of the squared differences between pixel values of the original and the reconstructed images. Lower MSE indicates better performance.

$$\text{MSE} = \frac{1}{m \cdot n} \sum_{i=1}^m \sum_{j=1}^n (I_{\text{original}}(i, j) - I_{\text{reconstructed}}(i, j))^2$$

Where:

$I_{\text{original}}(i, j)$ is the pixel value at position (i, j) in the original image.

$I_{\text{reconstructed}}(i, j)$ is the pixel value at position (i, j) in the reconstructed image.

m, n are the dimensions of the image.

3. Structural Similarity Index Measure (SSIM)

SSIM measures the perceptual similarity between two images, considering luminance, contrast, and structure. Higher SSIM indicates higher similarity.

$$\text{SSIM}(x, y) = \frac{(2\mu_x\mu_y + C_1)(2\sigma_{xy} + C_2)}{(\mu_x^2 + \mu_y^2 + C_1)(\sigma_x^2 + \sigma_y^2 + C_2)}$$

Where:

x and y are two image patches.

μ_x, μ_y = Mean of x and y .

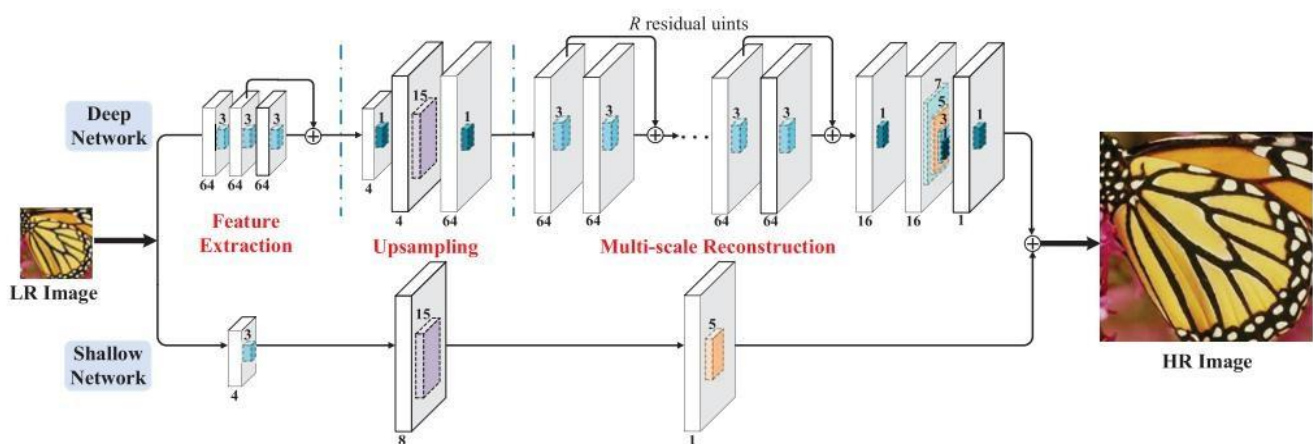
σ_x^2, σ_y^2 = Variance of x and y .

σ_{xy} = Covariance between x and y .

C_1 and C_2 = Stabilization constants to avoid division by zero.

Experimental Results

The proposed SRCNN model is evaluated on a set of degraded images to measure its performance before and after enhancement. The results are analyzed using PSNR, SSIM, and MSE metrics. The comparison between the degraded and reconstructed images is visualized using bar graphs for each metric.



Conclusion

In this paper, we implemented a Super-Resolution Convolutional Neural Network (SRCNN) to enhance the quality of low-resolution images. The model effectively learns to reconstruct high-resolution images from degraded inputs by leveraging a simple yet powerful three-layer architecture. The use of pre-trained weights significantly reduces training time while achieving satisfactory results. The experimental results demonstrate improvements in PSNR, SSIM, and MSE scores, proving the efficacy of the SRCNN model for image enhancement. The visual comparison using bar graphs provides further clarity on the improvement achieved through this approach.

References

1. Dong, C., Loy, C. C., He, K., & Tang, X. (2015). Image Super-Resolution Using Deep Convolutional Networks. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 38(2), 295-307.
2. Ledig, C., Theis, L., Huszár, F., Caballero, J., Aitken, A. P., Tejani, A., Totz, J., Wang, Z., & Shi, W. (2017). Photo-Realistic Single Image Super-Resolution Using a Generative Adversarial Network. *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, 4681-4690.
3. Wang, X., Yu, K., Wu, S., Gu, J., Liu, Y., Dong, C., Qiao, Y., & Change Loy, C. (2018). ESRGAN: Enhanced Super-Resolution Generative Adversarial Networks. *Proceedings of the European Conference on Computer Vision (ECCV) Workshops*.
4. Zhang, Y., Tian, Y., Kong, Y., Zhong, B., & Fu, Y. (2018). Residual Dense Network for Image Super-Resolution. *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, 2472-2481.
5. Simonyan, K., & Zisserman, A. (2015). Very Deep Convolutional Networks for Large-Scale Image Recognition. *arXiv preprint arXiv:1409.1556*.
6. Wang, Z., Bovik, A. C., Sheikh, H. R., & Simoncelli, E. P. (2004). Image Quality Assessment: From Error Visibility to Structural Similarity. *IEEE Transactions on Image Processing*, 13(4), 600-612.
7. Goodfellow, I., Pouget-Abadie, J., Mirza, M., Xu, B., Warde-Farley, D., Ozair, S., Courville, A., & Bengio, Y. (2014). Generative Adversarial Networks. *arXiv preprint arXiv:1406.2661*.
8. He, K., Zhang, X., Ren, S., & Sun, J. (2016). Deep Residual Learning for Image Recognition. *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, 770-778.

