

Super-Resolution of Compressed Images Using Residual Information Distillation Network

Yanqing Zhang ¹, Jie Li ^{1,2}, Nan Lin ¹, Yangjie Cao ¹ and Cong Yang ^{1,*}

¹ School of Cyber Science and Engineering, Zhengzhou University, Zhengzhou 450001, China

² Department of Computer Science and Engineering, Shanghai Jiao Tong University, Shanghai 200030, China

* Correspondence: wangyuanync@zzu.edu.cn

Abstract: Super-Resolution (SR) is a fundamental computer vision task, which reconstructs high-resolution images from low-resolution ones. Existing SR methods mainly recover images from clear low-resolution images, leading to unsatisfactory results when processing compressed low-resolution images. In the paper, we propose a two-stage SR method for compressed images, which consists of the Compression Artifact Removal Module (CARM) and Super-Resolution Module (SRM). The compressed low-resolution image is used to reconstruct the clear low-resolution image by CARM, and the high-resolution image is obtained by SRM. In addition, we propose a residual information distillation block to learn the texture details which are lost during the compression process. The proposed method has been validated and compared with the state of the art, and experimental results show that the proposed method outperforms other super-resolution methods in terms of visual effects and objective evaluation metrics.

Keywords: image super-resolution; compressed images; information distillation



Citation: Zhang, Y.; Li, J.; Lin, N.; Cao, Y.; Yang, C. Super-Resolution of Compressed Images Using Residual Information Distillation Network. *Electronics* **2023**, *12*, 1209. <https://doi.org/10.3390/electronics12051209>

Academic Editor: Byung-Gyu Kim

Received: 10 January 2023

Revised: 14 February 2023

Accepted: 23 February 2023

Published: 3 March 2023



Copyright: © 2023 by the authors. Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<https://creativecommons.org/licenses/by/4.0/>).

1. Introduction

Images are generally down-sampled and compressed so as to reduce space consumption and accelerate image transmission. The process of down-sampling causes loss of details in high-resolution images, and the subsequent compression also brings undesirable artifacts such as block effects. Reconstruction of images from down-sampled and compressed ones is therefore important for reducing storage consumption while retaining image details when viewing.

Deep learning (DL) has achieved encouraging results in various tasks, including image classification [1–5], object detection [6–10], and image segmentation [11–17]. Among them, super-resolution technology is currently a great topic, which refers to the generation of high-resolution (HR) images given the corresponding low-resolution (LR) ones. SR is now in urgent demand in many applications, such as intelligent surveillance [18], medical imaging [19,20], remote sensing [21,22], etc.

The key problem of compressed image super-resolution is to eliminate compression artifacts and preserve image details while increasing image resolution. The currently proposed deep learning-based super-resolution algorithms target uncompressed images, i.e., processed low-resolution images with only down-sampling or blurring degradation process, which leads to the fact that these super-resolution reconstruction algorithms for uncompressed images cannot effectively handle with compressed images. If super-resolution is performed directly on JPEG images it will aggravate the block effect and ringing effect, leading to poor visual effect [23]. Some algorithms have been proposed for compressed image reconstruction, but most of them are trained by decomposing this task into two independent subtasks and then using a joint level. Experimental results often show more pronounced compression artifacts and severe loss of details.

To address these issues, we propose a novel super-resolution algorithm to reconstruct compressed images. The model input uses three kinds of data: compressed low resolution

(C-LR) images as input, and LR and HR images as labels. We propose two modules: compression artifact removal module (CARM) and super-resolution module (SRM). To remove the compression artifacts in C-LR images, a two-stage joint loss function is used. The first of these stages use the LR image as supervised information, which greatly reduces the probability of image scaling errors in the super-resolution stage. To further make the generated images clearer, we propose a residual information-distillation module to learn more about the image features lost in the compression process. Finally, we use peak signal-to-noise ratio (PSNR) and structural similarity (SSIM) to evaluate the image quality, and the experimental results verify the effectiveness of the method. The main contributions of this paper are as follows:

Firstly, we design a compressed image super-resolution network consisting of two modules: Compression Artifact Removal Module (CARM) and Super-Resolution Module (SRM), using a joint loss function to train the network.

Secondly, we propose a novel residual information distillation block so as to efficiently learn the image features lost in the compression process.

Finally, experimental results show that the proposed model performs well and obtains better results on common evaluation metrics including peak signal-to-noise ratio (PSNR) than state-of-the-art.

2. Related Work

2.1. Problem Formulation

The traditional SR problem can be formulated as

$$Y = HX + n, \quad (1)$$

where X represents the HR image, H denotes the down-sampling and blur kernel, and n represents the additive noise. For super-resolution task for JPEG image, the downscaling process is shown in Figure 1. Compared with the traditional super-resolution problem, the JPEG compression low-resolution image degradation process adds JPEG compression. So (1) can be converted into (2):

$$Z = CY, \quad (2)$$

where Y is the low-resolution image defined by (1), C represents compression kernel, and Z denotes the compressed low-resolution image. This work mainly studies compression artifacts, so we ignore the additional noise. The degradation process of its low-resolution images as shown in Figure 1. Our goal is to generate the high-resolution image X from Z .

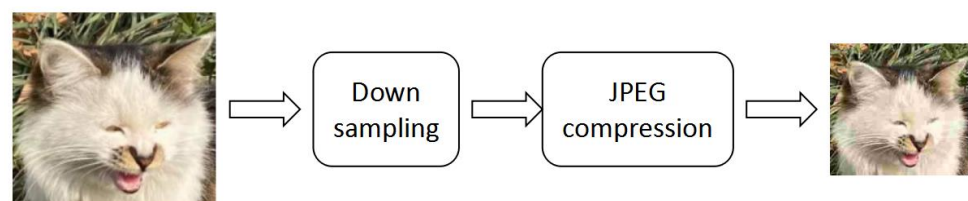


Figure 1. Degradation process for compressed LR images.

2.2. Image Super-Resolution

More SISR methods based deep learning have been proposed by researchers and achieved excellent results in image quality assessment metrics and visual effects. Dong et al. [24] originally proposed a three-layer convolutional super-resolution network SRCNN, which implemented feature block extraction and representation, nonlinear mapping, and image reconstruction. Subsequently, Kim et al. [25] alleviated the difficulty of network training by introducing residual learning, and the proposed 20-layer-deep model VDSR learns residual images instead of directly learning high-resolution images. To further fuse the shallow and deep features of the image, RDN [26] used dense connectivity to obtain richer reconstruction information and details. To address the situation that most

algorithms do not apply to mobile devices, Hui et al. [27] also proposed a lightweight information distillation network IMDN based on IDNs [28]. Overall, deep neural network-based super-resolution methods have good performance, but these algorithms only undergo a single down-sampling degradation process and are not applicable to directly solve the super-segmentation variability problem of compressed images. Therefore, the effectiveness is greatly reduced when applied to compressed images.

Several algorithms have been dedicated to addressing compression artifacts. AR-CNN [29] is based on a deep learning algorithm to address JPEG compression artifacts with four convolutional layers for feature extraction, feature enhancement, nonlinear mapping, and reconstruction in sequence. After that, the CISRDCNN [30] algorithm is invoked for the first time to solve the end-to-end super-resolution for compression image. It can effectively improve the image quality and reduce the compression artifacts. To preserve the functionality of each module and the relevance of the two subproblems, each module is first trained end-to-end individually, and finally, the whole network is trained by joint optimization.

However, compression artifacts removal and up-sampling in the existing SR method are considered two independent stages, which may result in more artifacts or excessive smoothing. Overall, the SR of compression images still needs to be improved.

3. Proposed Method

A novel super-resolution reconstruction network is proposed for compressed images in this work, and Figure 2 shows the detailed network framework.

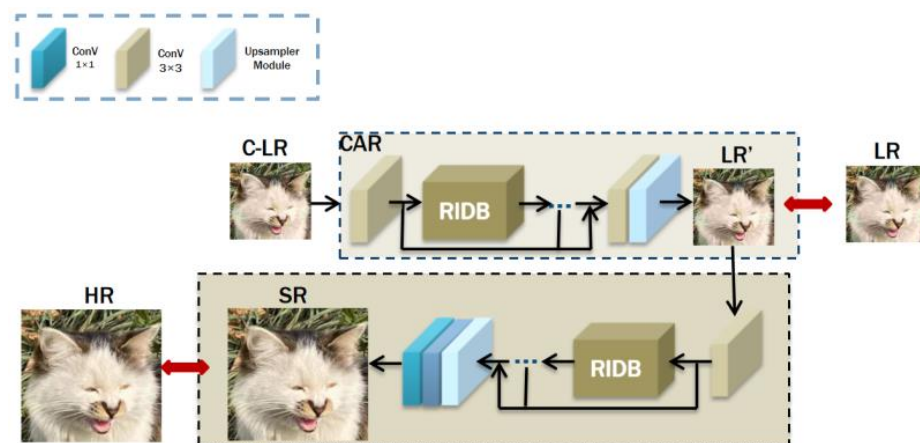


Figure 2. The overall pipeline of the proposed super-resolution reconstruction network.

3.1. General Framework

This paper aims to leverage compressed LR images to reconstruct HR images. During the training process, each sample includes three types of data: the C-LR image, the HR image, and the LR image. The C-LR image is processed by JPEG compression based on the LR image. Down-sampling is employed to generate the LR image. It is used as the ground truth in the first stage. In the second stage, the HR image is taken as the label image. The two stages of our method are presented separately as follows.

The method consists of two steps: Compression Artifact Removal Module (CARM) and Super-Resolution Module (SRM). Since the compression operation makes the low-resolution image lose more image details, reconstructing more image details based on the LR image is the key of the first stage. In Figure 2, the C-LR image is firstly reconstructed details lost during image compression by a compression artifact removal module, which consists of feature extraction, residual information distillation blocks, feature fusion layer, and reconstruction block. The main part of this stage is composed of multiple residual information distillation blocks (RIDBs), which are stacked to progressively refine the extracted features. The RIDBs will be described in Section 3.2. Finally, the part uses feature fusion layers as well as reconstruction blocks to reconstruct a clear low-resolution image.

The image super-resolution stage uses essentially the same network configuration as the first stage except for the final sub-pixel layer. Specifically, the second sub-module takes the clear LR image predicted by the first stage as input and outputs the HR image. This stage uses the same RIDB module as the first stage, the main reason is that each layer of the network in this module is good at learning the pixel-level feature representation. The process of low to high resolution image enlargement is achieved by sub-pixel layers, which will significantly reduce the training time since the sub-pixel layer only changes the image size at the last layer. The final model will output high-resolution images.

3.2. Residual Information Distillation Block

The proposed model is comprised of residual information distillation blocks, which are stacked to gradually refine the extracted features, where the RIDB structure is depicted in Figure 3.

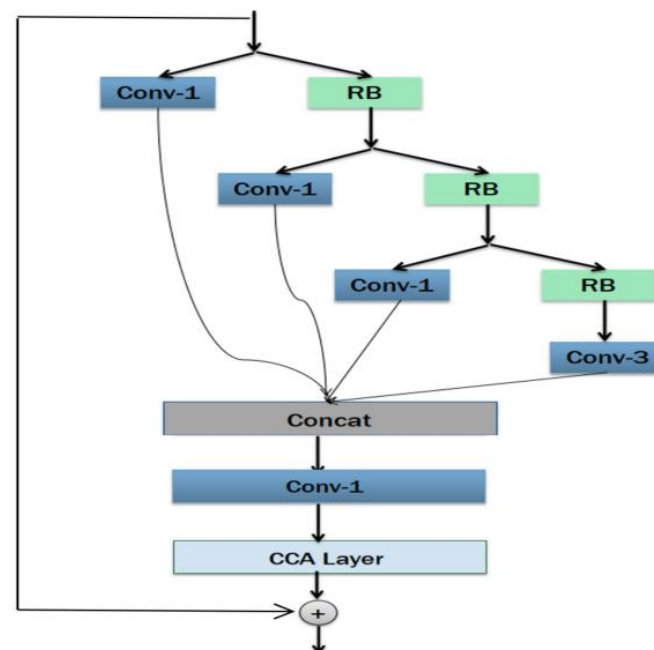


Figure 3. The architecture of RIDB.

The RIDB is composed of a residual block (RB) layer, a convolutional layer, and a contrast-aware channel attention (CCA) layer. The input features are split into two parts after a channel distillation operation: one part of the features is retained when the other part is sent to the next RIDB. The distillation operation compresses the feature channels in a fixed proportion, so 30% of the features are retained in this paper. The left-side retained features use a 1×1 convolution instead of a 3×3 convolution to decrease the number of parameters while remaining efficient. The split features on the right side enter the RB for deeper residual learning, where the RB consists of two 3×3 convolutions and an excitation module. Since RB is the body part of RIDB, 3×3 convolution can be used to capture the background information effectively and further refine the features. RB can learn deeper features from the residual learning without introducing any additional parameters.

3.3. Loss Functions

As illustrated in 3.1, the network is comprised of two stages: the CARM and the SRM. The stage 1 aims to reconstruct the C-LR image to the LR image by recovering the corresponding location pixel point in the C-LR and LR image. The objective of this stage can be expressed as making the C-LR image recover to the LR pixel value. Therefore, the L1

loss function is chosen in this stage which is prevalent in pixel-level tasks (image denoising, super-resolution, and image deblurring). The loss can be presented as (3):

$$L_{CAR} = \frac{1}{WH} \sum_{w=1}^W \sum_{h=1}^H \| I_{wh}^{LR} - F_{CAR}(I_{wh}^{C-LR}) \|, \quad (3)$$

where H and W respectively denote the height and the width of the C-LR image. F_{CAR} corresponds to the network of the first stage in (3).

The intermediate result of the previous stage is reconstructed by super-resolution in the second stage to change the resolution size of the image, such that a clear high-resolution image is obtained. To make the SRM generate more accurate SR, the loss function of stage II inherits the same loss of the previous excellent SRM, so the loss can be presented as Equation (4):

$$L_{SR} = \frac{1}{s^2WH} \sum_{w=1}^{sW} \sum_{h=1}^{sH} \| I_{sw \times sh}^{HR} - F_{SR}(I_{wh}^{L \ R'}) \|, \quad (4)$$

where s is the scaling factor. For the sake of generalizing most cases, the loss weights of both stages are set equal in proportion in the text, so the general loss function of this algorithm is shown as (5):

$$L_{total} = L_{CAR} + L_{SR}. \quad (5)$$

4. Experimental Results and Analysis

4.1. Dataset and Implementation Settings

The widely used dataset DIV2K in the image super-resolution field is employed in this work. The 1000 high-quality RGB images in DIV2K are divided into 800, 100, and 100 three parts for training, validation, and testing, respectively. In order to obtain the compressed LR image, the HR image is firstly down-sampled by bicubic to generate the uncompressed LR image, and the scale factor in the down-sampling process is set to 2 and 4. After that, the clear LR image is compressed by the JPEG encoder in MATLAB to obtain the compressed low-resolution image. The standard JPEG compression method is used in this experiment, and the compression quality factor QF is set to 20. The image dataset contains HR images, LR images, and C-LR images. The test data were selected from the widely used Set5 [31] and Set14 [32].

During training, the LR image is randomly cropped to 64×64 size as the model input, the training epochs is 1500 and the batch size is 32. The input image is randomly flipped horizontally or rotated by 90 degrees to enhance the data. This paper uses ADAM optimizer for optimization training, where $\beta_1 = 0.9$ and $\beta_2 = 0.999$, the learning rate is 2×10^{-4} initially and decreased half every 2×10^5 training rounds.

4.2. Evaluation Metrics

PSNR [33] and SSIM [34] are utilized to measure the effectiveness of the SR methods, where all values are calculated after converting the image of the RGB channel to the color space of the YCrCb channel for the Y channel is calculated. The full-reference image evaluation metric PSNR is widely leveraged in image restoration tasks such as SR and deblurring, which calculates the magnitude of the global pixel error between the original and reconstructed image to measure the image quality. The image similarity is evaluated by SSIM, which combines brightness, structure, and contrast. SSIM takes 0 to 1, where the high value represents the significant similarity.

4.3. Experimental Results

To verify the performance of the proposed method, this paper compares with the state-of-the-art works including Bicubic, ARCNN [29], VDSR [25], RCAN [26], IMDN [32], and CISRD CNN [31]. These algorithms are tested on Set5 and Set14 datasets using PSNR and SSIM as evaluation metrics. To ensure the fairness of the comparison experiments, the

models and codes for the comparison experiments are obtained from the URLs provided in the relevant papers, and the test datasets are compressed.

The experimental results of each algorithm when the down-sampling factors are 2 and 4 are presented in Table 1. In particular, the bold font indicates the optimal results for each item.

Table 1. The comparative results on PSNR and SSIM.

Model	Scale	Set5	Set14
		PSNR (dB)/SSIM	PSNR (dB)/SSIM
Bicubic	×2	27.434/0.786	24.131/0.642
ARCNN	×2	27.678/0.790	25.121/0.608
VDSR	×2	28.856/0.815	25.514/0.688
IMDN	×2	30.002/0.835	25.955/0.669
Ours	×2	30.953/0.872	26.269/0.704
Bicubic	×4	24.116/0.606	23.156/0.552
ARCNN	×4	24.548/0.641	23.534/0.537
VDSR	×4	24.802/0.624	23.825/0.586
IMDN	×4	25.474/0.701	24.202/0.603
Ours	×4	26.430/0.766	24.562/0.631

Obviously, the CNN delivers the best results on both PSNR and SSIM when the down-sampling factor is 2 and 4, with some improvement compared to Bicubic. VDSR and IMDN use the residual blocks to learn high-frequency information from the image and obtain better visual results while improving the network efficiency. In terms of PSNR and SSIM, the proposed method surpasses other methods on down-sampling factors of 2 and 4, meaning that our strategy is able to reduce noise well and recover more high-frequency details in more complex environments. In comparison with the second-best approach, our method can improve up to about 0.9 and 0.1 dB in PSNR and SSIM, with the largest gain coming from the set5 dataset, which is arguably the most frequently used set of data in these experiments.

In addition, many noise and block artifacts occurred in the images, which may contain clues to recover their textures and details. It also shows that methods with high performance can successfully remove noise and can distinguish between noise and artifacts, which can be used to retain some complex textures and high-frequency details.

In this paper, two images from the dataset Set5 are selected for comparison of the experimental results with a down-sampling factor of 2 and a QF of 20. Figure 4 represents the visual reconstruction effect of each algorithm on the same image. For better observation of comparative effects, a local area of the image is given in the figure for the comparison demonstration of various algorithms. Table 2 represents the PSNR and SSIM for each algorithm on these two images. It is clear from the figure and the table that Bicubic has a more noticeable compression micro-shadow and serious loss of image details than other deep learning-based algorithms, and the image lash details are difficult to observe. The reconstructed images of ARCNN and VDSR also have compression artifacts, the IMDN algorithm is capable of removing the majority of the compression artifacts, but the image becomes smooth. In contrast, the reconstructed images of the proposed method have the least compression macrophages and more details are preserved, which has the most effective visual effect.

Table 2. Comparison of the results on PSNR and SSIM in Baby and Bird.

Model PSNR (dB)/SSIM	Bicubic	ARCNN	VDSR	Ours
baby	30.030/0.801	30.532/0.853	31.224/0.841	32.133/0.899
bird	30.379/0.782	31.004/0.806	31.675/0.813	31.998/0.806

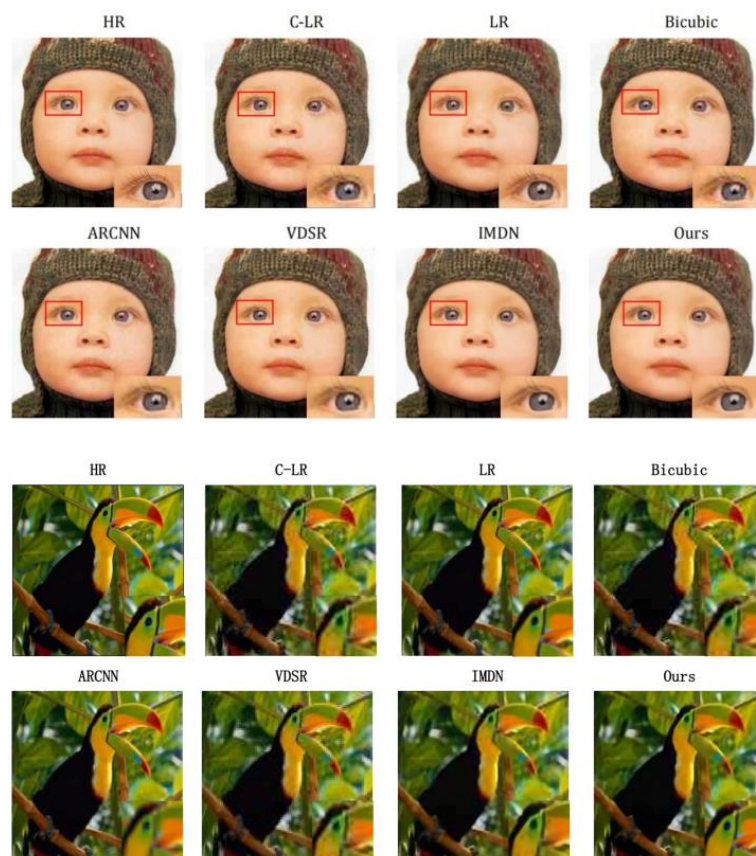


Figure 4. Visual comparison of different methods in Baby and Bird.

4.4. Ablation Study

The performance of the first stage CARM is tested here to demonstrate the effectiveness of our strategy.

In this section, the paper removes the CAR module from the model to verify its effectiveness. The C-LR image distortion is super-resolved directly, the network architecture of the SRM remains unchanged, and L1 loss is still employed for training. From Table 3, it can be seen that both PSNR and SSIM metrics are significantly improved after using the CARM. The reason is that, when the experiment is performed directly on the compressed image for super-resolution reconstruction, it leads to the micro-shadow and noise generated by the compression process is amplified in the super-resolution process, which results in undesirable visual effects. A single model cannot combine the recovery task and the SR task, so the CARM can perform supervised learning using LR to learn the mapping from C-LR to LR, thus tabulating better results in the reconstruction stage.

Table 3. Comparison of results for modules with and without CAR.

Model	Scale	Set5	Set14
		PSNR (dB)/SSIM	PSNR (dB)/SSIM
Model without CAR	$\times 2$	30.100/0.841	25.988/0.654
Model	$\times 2$	30.953/0.872	26.269/0.704
Model without CAR	$\times 4$	25.474/0.701	24.202/0.603
Model	$\times 4$	26.430/0.766	24.562/0.631

5. Conclusions

This paper proposes a two-stage SR reconstruction approach to address the reconstruction of HR images from compressed LR images. The network consists of CARM and SRM and incorporates residual information distillation blocks to extract hierarchical

features. Extensive experiments are conducted on real low-quality images and the results have shown that our method obtains better high-resolution images and better performance on objective evaluation metrics in comparison with the state-of-the-art. We demonstrate the application of the method in this paper to compressed images, allowing users to view clear images as well as facilitating post-visualization tasks of the images. The application of the algorithm can be extended to images and videos of other compression standards, such as jpeg2000 and HEVC. However, this task is still challenging because our method and the algorithms proposed so far still cannot accurately reconstruct the full texture of compressed images, which points to a direction for future research, which is to use generative adversarial networks to solve this problem.

Author Contributions: Conceptualization, J.L., Y.C. and N.L.; methodology, Y.Z. and C.Y.; software, Y.Z. and C.Y.; data curation, J.L. and Y.C.; writing—original draft preparation, Y.Z. and C.Y.; writing—review and editing, J.L., N.L., Y.C. and C.Y.; investigation, N.L. and Y.C.; supervision, J.L., N.L., Y.C. and C.Y.; validation, Y.Z. and J.L.; funding acquisition, Y.C. and J.L. All authors have read and agreed to the published version of the manuscript.

Funding: This research was funded by the Collaborative Innovation Major Project of Zhengzhou (grant number 20XTZX06013) and the National Natural Science Foundation of China (grant number 61972092).

Data Availability Statement: Not applicable.

Conflicts of Interest: The authors declare no conflict of interest.

References

1. Yang, X.; Wu, W.; Liu, K.; Kim, P.W.; Sangaiah, A.K.; Jeon, G. Long-Distance Object Recognition with Image Super Resolution: A Comparative Study. *IEEE Access* **2018**, *6*, 13429–13438. [\[CrossRef\]](#)
2. Siadari, T.S.; Han, M.; Yoon, H. GSR-MAR: Global Super-Resolution for Person Multi-Attribute Recognition. In Proceedings of the International Conference on Computer Vision Workshop (ICCVW), Seoul, Republic of Korea, 27–28 October 2019; pp. 1098–1103.
3. Que, Y.; Dai, Y.; Ji, X.; Leung, A.K.; Chen, Z.; Jiang, Z.; Tang, Y. Automatic classification of asphalt pavement cracks using a novel integrated generative adversarial networks and improved VGG model. *Eng. Struct.* **2023**, *277*, 115406. [\[CrossRef\]](#)
4. Li, S.; Song, W.; Fang, L.; Chen, Y.; Ghamisi, P.; Benediktsson, J.A. Deep Learning for Hyperspectral Image Classification: An Overview. *IEEE Trans. Geosci. Remote Sens.* **2019**, *57*, 6690–6709. [\[CrossRef\]](#)
5. Chen, C.F.R.; Fan, Q.; Panda, R. Crossvit: Cross-attention multi-scale vision transformer for image classification. In Proceedings of the IEEE/CVF International Conference on Computer Vision, Montreal, BC, Canada, 11–17 October 2021; pp. 357–366. [\[CrossRef\]](#)
6. Wu, F.; Duan, J.; Ai, P.; Chen, Z.; Yang, Z.; Zou, X. Rachis detection and three-dimensional localization of cut off point for vision-based banana robot. *Comput. Electron. Agric.* **2022**, *198*, 107079. [\[CrossRef\]](#)
7. Bochkovskiy, A.; Wang, C.Y.; Liao, H.Y.M. Yolov4: Optimal Speed and Accuracy of Object Detection. *arXiv* **2020**, arXiv:2004.10934.
8. Ancha, S.; Nan, J.; Held, D. Combining deep learning and verification for precise object instance detection. *arXiv* **2019**, arXiv:1912.12270.
9. Pang, Y.; Cao, J.; Wang, J.; Han, J. JCS-Net: Joint Classification and Super-Resolution Network for Small-Scale Pedestrian Detection in Surveillance Images. *IEEE Trans. Inf. Secur.* **2019**, *14*, 3322–3331. [\[CrossRef\]](#)
10. Zhang, Y.; Bai, Y.; Ding, M.; Xu, S.; Ghanem, B. KGSnet: Key-point-guided super-resolution network for pedestrian detection in the wild. *IEEE Trans. Neural Netw. Learn. Syst.* **2020**, *32*, 2251–2265. [\[CrossRef\]](#) [\[PubMed\]](#)
11. Minaee, S.; Boykov, Y.Y.; Porikli, F.; Plaza, A.J.; Kehtarnavaz, N.; Terzopoulos, D. Image segmentation using deep learning: A survey. *IEEE Trans. Pattern Anal. Mach. Intell.* **2021**, *44*, 3523–3542. [\[CrossRef\]](#) [\[PubMed\]](#)
12. Hatamizadeh, A.; Tang, Y.; Nath, V.; Yang, D.; Myronenko, A.; Landman, B.; Xu, D. Unetr: Transformers for 3d medical image segmentation. In Proceedings of the IEEE/CVF Winter Conference on Applications of Computer Vision, Waikoloa, HI, USA, 3–8 January 2022; pp. 1748–1758. [\[CrossRef\]](#)
13. Dai, D.; Wang, Y.; Chen, Y.; Van Gool, L. Is image super-resolution helpful for other vision tasks? In Proceedings of the 2016 IEEE Winter Conference on Applications of Computer Vision (WACV), Lake Placid, NY, USA, 7–10 March 2016; pp. 1–9. [\[CrossRef\]](#)
14. Guo, Z.; Wu, G.; Song, X.; Yuan, W.; Chen, Q.; Zhang, H.; Shi, X.; Xu, M.; Xu, Y.; Shibasaki, R.; et al. Super-Resolution Integrated Building Semantic Segmentation for Multi-Source Remote Sensing Imagery. *IEEE Access* **2019**, *7*, 99381–99397. [\[CrossRef\]](#)
15. Wang, L.; Li, D.; Zhu, Y.; Tian, L.; Shan, Y. Dual Super-Resolution Learning for Semantic Segmentation. In Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition, Seattle, WA, USA, 13–19 June 2020; pp. 3773–3782. [\[CrossRef\]](#)

16. Bhojanapalli, S.; Chakrabarti, A.; Glasner, D.; Li, D.; Unterthiner, T.; Veit, A. Understanding robustness of transformers for image classification. In Proceedings of the IEEE/CVF International Conference on Computer Vision, Montreal, BC, Canada, 11–17 October 2021; pp. 10231–10241.
17. Wang, P.; Fan, E.; Wang, P. Comparative analysis of image classification algorithms based on traditional machine learning and deep learning. *Pattern Recognit. Lett.* **2021**, *141*, 61–67. [[CrossRef](#)]
18. Zhang, L.; Zhang, H.; Shen, H.; Li, P. A super-resolution reconstruction algorithm for surveillance images. *Signal Process.* **2010**, *90*, 848–859. [[CrossRef](#)]
19. Huang, Y.; Shao, L.; Frangi, A.F. Simultaneous Super-Resolution and Cross-Modality Synthesis of 3D Medical Images Using Weakly-Supervised Joint Convolutional Sparse Coding. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, Honolulu, HI, USA, 21–26 July 2017; pp. 5787–5796. [[CrossRef](#)]
20. Greenspan, H. Super-Resolution in Medical Imaging. *Comput. J.* **2008**, *52*, 43–63. [[CrossRef](#)]
21. Lei, S.; Shi, Z.; Zou, Z. Super-Resolution for Remote Sensing Images via Local–Global Combined Network. *IEEE Geosci. Remote Sens. Lett.* **2017**, *14*, 1243–1247. [[CrossRef](#)]
22. Zhang, S.; Yuan, Q.; Li, J.; Sun, J.; Zhang, X. Scene- adaptive remote sensing image super-resolution using a multiscale attention network. *IEEE Trans. Geosci. Remote Sens.* **2020**, *58*, 4764–4779. [[CrossRef](#)]
23. Dong, C.; Loy, C.C.; He, K.; Tang, X. Learning a Deep Convolutional Network for Image Super-Resolution. In Proceedings of the ECCV 2014, Zurich, Switzerland, 6–12 September 2014. [[CrossRef](#)]
24. Dong, C.; Loy, C.C.; Tang, X. Accelerating the Super-Resolution Convolutional Neural Network. In *Computer Vision-Eccv 2016*; Li, P., Leibe, B., Eds.; Springer International Publishing: New York, NY, USA, 2016; pp. 391–407.
25. Kim, J.; Lee, J.K.; Lee, K.M. Accurate image super-resolution using very deep convolutional networks. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, Las Vegas, NV, USA, 27–30 June 2016; pp. 1646–1654.
26. Xiong, Z.; Sun, X.; Wu, F. Robust web image/video super-resolution. *IEEE Trans. Image Process.* **2010**, *19*, 2017–2028. [[CrossRef](#)] [[PubMed](#)]
27. Hui, Z.; Gao, X.; Yang, Y.; Wang, X. Lightweight image super-resolution with information multi-distillation network. In Proceedings of the 27th ACM International Conference on Multimedia, Nice, France, 21–25 October 2019; pp. 2024–2032.
28. Hui, Z.; Wang, X.; Gao, X. Fast and accurate single image super-resolution via information distillation network. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, Salt Lake City, UT, USA, 18–22 June 2018; pp. 723–731.
29. Yu, K.; Dong, C.; Loy, C.C.; Tang, X. Deep convolution networks for compression artifacts reduction. *arXiv* **2016**, arXiv:1608.02778.
30. Chen, H.; He, X.; Ren, C.; Qing, L.; Teng, Q. CISRDCNN: Super-resolution of compressed images using deep convolutional neural networks. *Neurocomputing* **2018**, *285*, 204–219. [[CrossRef](#)]
31. Bevilacqua, M.; Roumy, A.; Guillemot, C.; Alberi Morel, M.-L. Low-Complexity Single-Image SUPER-RESolution based on Nonnegative Neighbor Embedding. In *British Machine Vision Conference BMVA*; BMVA Press: Durham, UK, 2012.
32. Zeyde, R.; Elad, M.; Protter, M. On single image scale-up using sparse-representations. In Proceedings of the 2010 International Conference on Curves&Surfaces, Avignon, France, 24–30 June 2010; Springer: Berlin, Germany, 2010; pp. 711–730.
33. Sheikh, H.R.; Sabir, M.F.; Bovik, A.C. A Statistical Evaluation of Recent Full Reference Image Quality Assessment Algorithms. *IEEE Trans. Image Process.* **2006**, *15*, 3440–3451. [[CrossRef](#)] [[PubMed](#)]
34. Wang, Z.; Bovik, A.C.; Sheikh, H.R.; Simoncelli, E.P. Image Quality Assessment: From Error Visibility to Structural Similarity. *IEEE Trans. Image Process.* **2004**, *13*, 600–612. [[CrossRef](#)] [[PubMed](#)]

Disclaimer/Publisher’s Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.