

Low-Resolution Face Recognition: A Review of Methods and Data

Yesun Utomo^{1*}, Kelvin²

^{1,2}Computer Science Department, School of Computer Science,
Bina Nusantara University,
Jakarta, Indonesia 11480
yesun.utomo@binus.edu; kelvin007@binus.ac.id

*Correspondence: yesun.utomo@binus.edu

Abstract – This paper provides a review of previous studies in Low-Resolution Face Recognition (LRFR), specifically focusing on cross-resolution Face Recognition (FR) methods. While state-of-the-art deep learning FR systems achieve high accuracy on high-resolution (HR) images, they are generally unsuitable for low-resolution (LR) images frequently encountered in applications like surveillance systems, where faces often have low pixel counts due to capture conditions. Cross-resolution FR, which compares an HR image with an LR image, presents a significant challenge due to the distinct visual properties of images at different resolutions. The paper discusses two primary approaches to address this problem: super-resolution (SR), which is a transformative method that aims to construct HR images from LR ones, and unified feature space (UFS), a non-transformative method that maps facial features from varying resolutions into a shared feature space. This work summarizes both SR and UFS methods. Based on the review, the paper concludes that non-transformative (UFS) methods are more suitable for future directions. This recommendation is driven by their lower computational power requirements, proven effectiveness in real-world implementations such as mobile devices and drones, and alignment with current technological trends. The paper also emphasizes the need for further research using real or natural LR face images to identify degradation patterns and compare results between real and artificially generated LR images.

Keywords: Deep Learning; Face Recognition; Low-Resolution Face Recognition

I. INTRODUCTION

In recent years, Face Recognition (FR) systems have seen significant advancements in deep learning architectures. Some state-of-the-art models have achieved more than 99% accuracy on the LFW dataset. However, these models require high-resolution (HR) images, including pre-processing steps like face alignment. As such, they are unsuitable for low-resolution (LR) images. LR images are typically found in surveillance systems which often produce low-pixel-counts face images.

Some previous work, such as that by (Knoche et al., 2021) that uses ResNet50 has shown that image resolutions below 50 x 50 px significantly reduce FR accuracy. Other work that uses IResNet (Duta et al., 2020) and MobileNet (Sandler et al., 2018) also gives similar reports.

In general, there are two FR scenarios depending on the image resolution (Knoche et al., 2023). First is FR by comparing 2 images where both are LR. The second one is FR by comparing 2 images where one is HR and the other is LR, also known as cross-resolution FR method. The latter is harder because of the distinct visual properties of HR and LR images. There are a few approaches to solve cross-resolution FR issues:

- 1) Super-resolution (SR) is a method that uses LR images to construct HR images, often called transformative approaches.
- 2) Map or project face features from existing LR and HR into similar space to learn the Unified Feature Space (UFS), often called non-transformative approach.

In this work, we review previous studies on Low Resolution Face Recognition (LRFR),

focusing specifically on cross-resolution FR methods, comparing SR and UFS approaches. This paper gives the following contributions:

- Summary of Super Resolution methods.
- Summary of Unified Feature Space methods.
- Conclude which method is better for future directions.

One application of LRFR is in surveillance systems. Faces in these systems are often captured at considerable distances, high angles, and under poor lighting conditions. As a result, detected faces often provide low feature information. Although HR face recognitions have almost 100% accuracy, LRFR is still quite a challenge. In this section, we will discuss and analyze previous works conducted using SR / transformation-based approach and UFS / non-transformation-based approach).

II. METHODS

This paper selected relevant papers related to low-resolution face recognition frameworks. We employed a systematic approach using Google Scholar. Keywords such as “low-resolution face recognition” and “low resolution face recognition dataset” are utilized to identify pertinent literature. Papers selected are no earlier than 2016 to preserve relevancy and to observe changes of LRFR approaches over the years. The search process includes screening titles, abstracts, and keywords to ensure relevancy. Citation tracking and reference lists are also used to collect more relevant studies. After the screening process, selected papers will be examined thoroughly to assess their adherence to the LRFR methodologies. Papers that met the criteria are included for further analysis, while those that did not meet the criteria are excluded.

Once the final set of papers are gathered, we perform a comprehensive review of the methodologies used in each paper. This involved extracting key details about the research designs, data collection, preprocessing techniques, and modeling approaches for LRFR.

The analysis encompasses the summary and comparison of Unified Feature Space & Super Resolution derivation such as: Prior-Guided, Identity Preserving, and Reference. This also includes relevant datasets that can be used to

support the two. The findings are organized and presented in a structured manner to provide a clear understanding of the employed methodologies.

III. RESULTS AND DISCUSSION

In this section, we will present papers that contribute to LRFR. To organize the contributions effectively, the papers are arranged in chronological order based on the year of publication. This allows us to track research progression and observe the evolution of the field. Datasets that can be used to LRFR will be presented in this section as well.

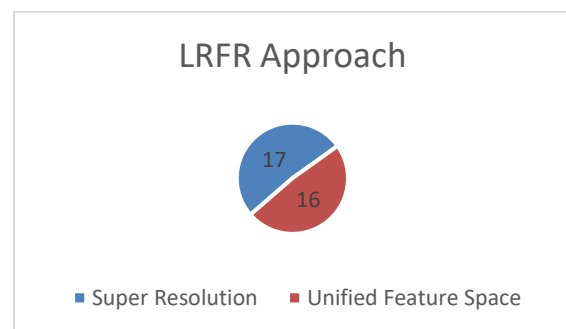


Figure 1. Distribution of LRFR Approach

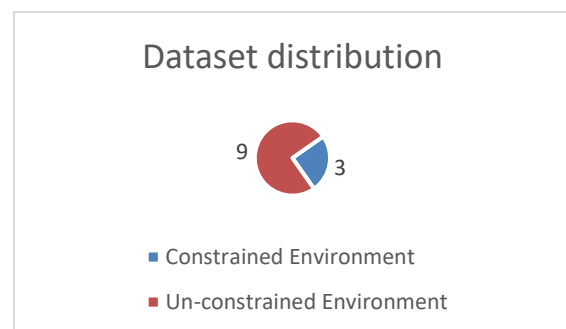


Figure 2. Distribution of Dataset

One application of LRFR is in surveillance systems. Faces in these systems are often captured at considerable distances, high angles, and under poor lighting conditions. As a result, detected faces often provide low feature information. Although HR face recognitions have almost 100% accuracy, LRFR is still quite a challenge. In this section, we will discuss and analyze previous works conducted using SR / transformation-based approach and UFS / non-transformation-based approach).

3.1 Super Resolution (transformation-based approach)

Super Resolution works by transforming LR face images into HR face images, which are then used for FR. Super Resolution itself can be

categorized into several “branch” depending on the approach. For example: General Face SR, Prior-guided Face SR, Attribute-constrained Face SR, Identity-preserving Face SR, and Reference Face SR, etc. (Jiang et al., 2021).

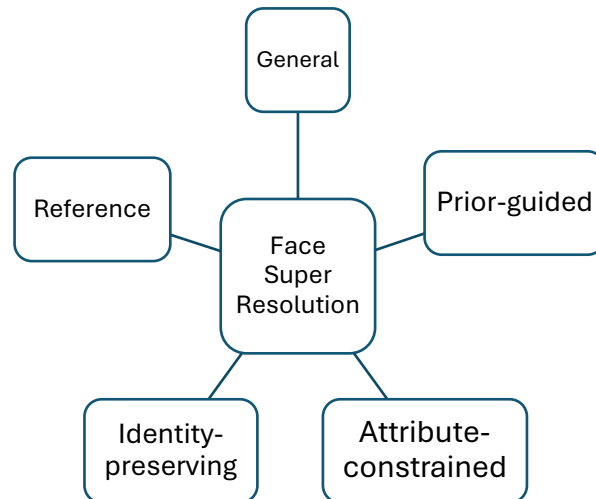


Figure 3. Taxonomy of face super-resolution

Each with their own pros and cons. For example, General Face SR methods aim to create efficient networks. However, they ignore prior information of human faces which results in HR images with mediocre facial structures. Among the 5 mentioned branches, we will discuss Prior-guided, Identity-preserving, and Reference. The reason is because these 3 are among the most popular methods.

3.1.1 Prior-guided

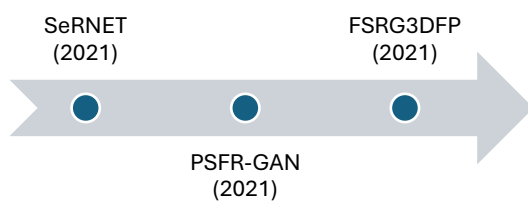


Figure 4. Milestones of pre-prior FR method

Prior-guided Face SR works by extracting facial prior information and using that information for HR reconstruction. This method derives its name from the use of prior information, such as facial parsing maps, heatmaps, and landmarks which guide the reconstruction process. This aims to create HR face images with better facial structures. This Face SR method can be divided further into

several approaches. SeRNet (Yu et al., 2021) pre-trains a face structure extraction network that will be used to extract facial parsing maps. PSFR-GAN (C. Chen et al., 2021) to calculate matrix loss of each semantic regions, designs a new semantic aware style loss. FSRG3DFP (Hu et al., 2021) uses a different approach by using 3D facial priors instead of the usual 2D priors. This method learns facial details and facial components in 3D by the spatial feature transform block. Those mentioned methods are called pre-prior because they extract prior estimation before performing super resolution, thus ignoring the correlation between face structure prior estimation and SR enhancement.

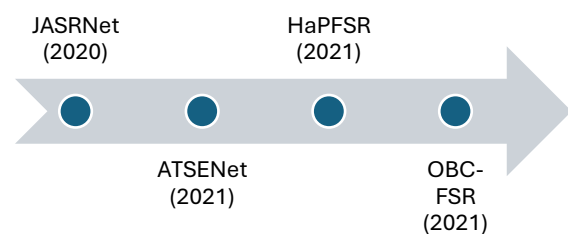


Figure 5. Milestones of parallel-guided FR method

Other prior-guided methods work by performing prior estimation and super-resolution in parallel. More specifically, they create prior estimation network and super-

resolution network, train them in parallel and take ground truth of facial prior information to calculate prior-based loss. JASRNet (Yin et al., 2020) uses a shared encoder which performs the 2 jobs simultaneously. ATSENet (M. Li et al., 2021) also could perform 2 jobs simultaneously. But unlike JASRNet, it feeds features from prior estimation network branch into a unit in the super-resolution branch called feature fusion unit. Other methods like HaPFSR (C. Wang et al., 2021) and OBC-FSR (J. Li et al., 2021) also works similarly.

3.1.2 Identity Preserving

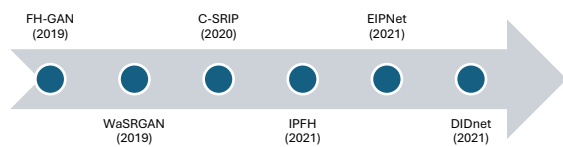


Figure 6. Milestones of identity-preserving FR method

Identity-preserving Face SR works by maintaining identity consistency of LR face image and generated HR face images with the purpose of performance improvement of downstream FR. Identity-preserving can be divided into 2 methods: *FR-based* or *Pairwise Data-based*. The former uses FR network to find identity loss. There are 2 main components: super-resolution model, a pre-trained face recognition network, and optional discriminators. The super-resolution model resolves the input of LR face image, generating identity value which is fed into the pre-trained FR network. The recent previous works are FH-GAN (face hallucination generative adversarial network) (Bayramli et al., 2019), WaSRGAN (H. Huang et al., 2019), IPFH (Identity preserving face hallucination) (F. Cheng et al., 2021), and C-SRIP (Grm et al., 2020). Another work called EIPNet (Kim et al., 2020) contributes by minimizing distortions created in SR by using lightweight edge block and identity information. The edge block is lightweight by only using 1 convolutional layer and 1 pooling layer. Extracted edge information from images are concatenated to the original feature maps in multiple scales. DIDnet (F. Cheng et al., 2021) contributes by integrating constraints to 2 closed loops network. The first loop network is

responsible for generating HR face images from LR ones, where the constraint preserves the identity in HR features space. The second loop network is responsible for learning image degradation process.

3.1.3 Reference

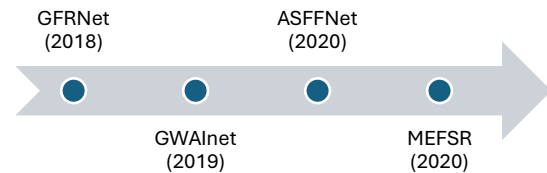


Figure 7. Milestones of reference FR method

Unlike previous methods that use only LR face image as input, *Reference* also uses HR face images of the same identity that comes within provided dataset. These HR face images give face details which are used for *reference* for face reconstruction; thus the method name comes from. There are 2 different approaches in reference: single reference image or multiple reference images.

In single reference image approach, as the name suggests, only uses one HR face image that shares the same identity with a LR face image which serves as reference. However, the single reference image may have different poses or facial expressions that could hamper face reconstruction. To fix this problem, a face alignment needs to be performed between the reference and the LR image. Then both the reference and the LR image are fed to the reconstruction network to perform SR. (X. Li et al., 2018) uses single reference image method and perform face alignment of the reference and LR image using face landmarks. Then they fed those images into the reconstruction network. (Dogan et al., 2019) uses similar approaches but perform face alignment using flow field. After the alignment process, features are extracted from the reference image and fed into the reconstruction network.

As for multiple reference images, (X. Li et al., 2020) is the first to explore this method called ASFFNet. Among the reference images, they select the image that has the most similar pose and expression to the LR image using guidance selection module. Any existing differences in alignment or illumination between reference image and LR image are mitigated using weighted least-square

alignment (Schaefer et al., 2006) and AdaIN (X. Huang & Belongie, 2017). Other work by (K. Wang et al., 2020) called MEFSR uses all reference images and fed them into weighted pixel average module for feature extraction and face reconstruction.

3.2 Unified Feature Space (Non-transformation-based Approach)

Non-transformation approach works by directly projecting facial features from various resolution images into a common feature space. Most of these methods use CNN like (Zangeneh et al., 2020) combined with a non-linear coupled mapping architecture. (Massoli et al., 2020) uses a teacher-student method. (Lu et al., 2018) uses deep coupled ResNet model

consisting of 1 trunk network to extract facial features and 2 branches of network transform HR and corresponding LR features into common feature space. (Talreja et al., 2019) proposed a coupled GAN and multiple loss functions utilized in attribute-guided cross-resolution FR model. (Ge et al., 2018) contributes a low computational method and uses a new learning approach with selective knowledge distillation. While (Sun et al., 2020) uses a shared classifier between HR and LR images to narrow domain gap. This work embeds multiple hierarchy loss into the intermediate layers to reduce the distance of intermediate features and avoids the over usage of intermediate features.



Figure 8. Milestones of non-transformational FR method

(Knoche et al., 2021) uses BT-M model. It was trained with half the number of LR images in each batch. In addition of BT-M model, they also use ST-M1 and ST-M2 where both incorporate a Siamese network structure that enables optimization in respect to additional feature distance loss. Another work by (Zeng et al., 2016) also uses a similar BT-M model, but they use a resolution-invariant deep network train it directly with unified LR and HR images. (Mudunuri et al., 2018) uses cross-resolution contrastive loss on features of 2 separate network branches. One branch focuses on HR images and the other focuses on LR images. (Lai & Lam, 2021) solved the problem of deep

Siamese network structure and combined cross-resolution triplet loss with a classification loss. While (Zha & Chao, 2019) also applied cross-resolution triplet loss like (Lai & Lam, 2021) combined with 2 branch networks like (Mudunuri et al., 2018). (Terhörst et al., 2021) aims to complete a distinct goal by focusing on image quality, pose, and age variations unlike previous works that focused only on image quality. This is achieved through combining quality-aware comparison score, model-specific face image quality, and FR model based on magnitude-aware angular margin loss. While (Zhao, 2021) uses a new technique for correlation feature-based FR, an unusual

method among other works. Recent work by (Knoche et al., 2023) uses Octuplet Loss to allow any network directly to learn the connection between HR and LR images while maintaining HR performance. Octuplet Loss is 4 different triplet loss inspired by (Lai & Lam, 2021) and (Gómez-Silva et al., 2020), however unlike (Lai & Lam, 2021), (Knoche et al., 2023) uses fine-tuning instead of training the model from scratch.

Recent technology of Unmanned Aerial Vehicles (UAVs) or drones also provides potential to perform LRFR. Due to the huge distance between the drone and face target, combined with limited battery capacity of the drone, it necessitated LRFR models that have both accuracy and lightweight computation. (Alansari et al., 2024) uses EfficientFaceV2S to perform LRFR on drones using drone-captured face image dataset. This method uses pretrained model from EfficientNetV2S and modified it by slowing down the sampling of the first convolution layer to 1 stride from the original of 2 strides. Accuracy and lightweight computation necessities also apply to mobile devices when performing FR in authentication. (Alansari et al., 2023) uses GhostFaceNets that uses PreLu activation function to improve model discrimination. This method uses Ghost bottleneck layers and some convolutions to only use 1x512 dimensions of embeddings.

3.3 Dataset

A high-quality dataset is essential for optimal results. In this case the dataset must contain labeled faces of people, either in LR or HR. Some datasets also contain extra information like gender, hair color, facial landmarks, facial heatmaps, and occlusions. In this section, we divide the datasets into 2 groups: dataset that contains LR and HR images (unconstrained environment), and dataset that contains only HR images (constrained environment). Optionally, researchers can generate their own dataset by creating LR images from existing HR datasets. Unconstrained environment datasets also provide various poses, illumination, occlusions, or small face pixels due to the huge distance between the subject and the camera. While constrained environment datasets also provide those variations, it is lackluster compared to the former.

3.3.1 Dataset in unconstrained environment

Table 1. Summary of public datasets in unconstrained environments.

Dataset	Image Number	Identities
AR Face	4.016	116
MegaFace Challenge 2	4.7 million	672.000
Youtube Faces	3.425	1595
SCface	4.160	130
UCCSface	70.000	130
LFW	13.233	5749
CASIA-WebFace	500.000	10.000
DroneFace	66	11
TinyFace	169.403	5139

- 1) AR Face: The AR Face dataset (Benavente, 1998) contains 116 identities (63 men and 63 women) with different facial expressions, illumination, and occlusions like sunglasses or scarves under strictly controlled conditions.
- 2) MegaFace Challenge 2: The MegaFace challenge 2 dataset (Nech & Kemelmacher-Shlizerman, 2017) contains 4.7 million face images and more than 672.000 identities. Subset selection can be performed on this dataset with tight bounding boxes to generate an LR subset of this dataset that contains faces smaller than 50x50 pixels (P. Li et al., 2018).
- 3) YouTube Faces: The YouTube Faces Database (Wolf et al., 2011) contains 3425 videos of 1595 identities. This dataset was designed to study the problem of unconstrained FR in videos.
- 4) SCface: The SCface dataset (Grgic et al., 2011) contains 4160 static face images (in the visible and infrared spectra) of 130 identities collected in an uncontrolled indoor environment using 5 video surveillance cameras of various qualities.
- 5) UCCSface: The UnConstrained College Students (UCCS) dataset (Sapkota & Boulton, 2013) Contains 180 identities of people walking on University of Colorado sidewalk from 100 to 150 meters, at 1 frame per second. The dataset consists of more than 70.000 hand-cropped face regions. Although the data are captured by HD

cameras, the face regions are tiny due to the large distance and contain a lot of noise or blurriness.

- 6) LFW: Labeled Faces in the Wild (LFW) dataset (G. B. Huang et al., 2008) contains 13233 face images of more than 5000 identities taken from public videos, pictures, etc. It is noted that faces other than the labeled target should be treated as background. 4069 identities only have a single image in the dataset while the rest have 2 or more images.
- 7) CASIA-WebFace: CASIA-WebFace dataset contains around 500000 images of 10000 identities. It was created to expand the LFW dataset.
- 8) DroneFace: DroneFace dataset (Hsu & Chen, 2017) contains 66 images of 11 identities (7 males, 4 females). The face images are taken by drones from various combinations of heights and distances. It is used for FR onto drones.
- 9) TinyFace: TinyFace dataset (Z. Cheng et al., 2018) contains more than 160000 images of 5000 identities. The images are captured in an open environment, thus, have real LR images.

3.3.2 Dataset in constrained environment

Table 2. Summary of public datasets in Constrained Environment

Dataset	Image Number	Identities
CACD200	160.000	2.000
VGGFace2	3.31 million	9.131
UMDFaces	367.888	8.277

- 1) CACD200: CACD200 dataset (B.-C. Chen et al., 2015) contains more than 160000 images of 2000 identities. The identities are celebrities with ages ranging from 16 to 62. This dataset can also be used to perform FR across ages.
- 2) VGGFace2: VGGFace2 dataset (Cao et al., 2017) contains more than 3 million images of more than 9000 identities. Each identity has an average of 360 images. The images are taken from Google Image Search and contain many variations of age, pose, ethnicity, illumination, and professions.

- 3) UMDFaces: UMDFaces dataset (Bansal et al., 2016) contains around 360000 images of more than 8000 identities. The images contain face key points, genders, and bounding boxes for faces.

IV. CONCLUSION

In this paper, we discussed two primary methods of LRFR: the transformation method and the non-transformation method. The transformation method or Super Resolution method can be divided into 5 sub-methods. We have discussed 3 of these sub-methods: prior-guided, identity-preserving, and reference because they are most popular. All 3 use large training models and require high computation power to reconstruct HR images which are still difficult to implement in the real world. As for non-transformation methods, it requires much less computing power and is much easier to implement in the real world as proven by several works in mobile devices (Alansari et al., 2023) and drones (Alansari et al., 2024) that have limited battery power. For future work, we suggest using non-transformative methods because they require less computing power, produce satisfying results proven by previous works, and current tech trends are mobile devices and drones.

Further research into using real or natural LR face images is also necessary to find patterns in the degradation process between LR and HR face images of the same person. This also includes comparing which will yield better results in LRFR between using real LR and artificial LR face images.

REFERENCES

- Alansari, M., Alnuaimi, K., Alansari, S., Javed, S., Shoufan, A., Zweiri, Y., & Werghi, N. (2024). *EfficientFaceV2S: A Lightweight Model and a Benchmarking Approach for Drone-Captured Face Recognition*. <https://ssrn.com/abstract=4698437>
- Alansari, M., Hay, O. A., Javed, S., Shoufan, A., Zweiri, Y., & Werghi, N. (2023). GhostFaceNets: Lightweight Face Recognition Model from Cheap Operations. *IEEE Access*, 11, 35429–35446.

- <https://doi.org/10.1109/ACCESS.2023.3266068>
- Bansal, A., Nanduri, A., Castillo, C., Ranjan, R., & Chellappa, R. (2016). *UMDFaces: An Annotated Face Dataset for Training Deep Networks*. <http://arxiv.org/abs/1611.01484>
- Bayramli, B., Ali, U., Qi, T., & Lu, H. (2019). *FH-GAN: Face Hallucination and Recognition using Generative Adversarial Network*. <http://arxiv.org/abs/1905.06537>
- Benavente, R. (1998). *The AR face database*. <https://www.researchgate.net/publication/243651904>
- Cao, Q., Shen, L., Xie, W., Parkhi, O. M., & Zisserman, A. (2017). *VGGFace2: A dataset for recognising faces across pose and age*. <http://arxiv.org/abs/1710.08092>
- Chen, B.-C., Chen, C.-S., & Hsu, W. H. (2015). Face Recognition and Retrieval Using Cross-Age Reference Coding With Cross-Age Celebrity Dataset. *IEEE Transactions on Multimedia*, 17(6), 804–815. <https://doi.org/10.1109/TMM.2015.2420374>
- Chen, C., Li, X., Yang, L., Lin, X., Zhang, L., & Wong, K.-Y. K. (2021). *Progressive Semantic-Aware Style Transformation for Blind Face Restoration*. <http://arxiv.org/abs/2009.08709>
- Cheng, F., Lu, T., Wang, Y., & Zhang, Y. (2021). Face Super-Resolution Through Dual-Identity Constraint. *2021 IEEE International Conference on Multimedia and Expo (ICME)*, 1–6. <https://doi.org/10.1109/ICME51207.2021.9428360>
- Cheng, Z., Zhu, X., & Gong, S. (2018). *Low-Resolution Face Recognition*. <http://arxiv.org/abs/1811.08965>
- Dogan, B., Gu, S., & Timofte, R. (2019). *Exemplar Guided Face Image Super-Resolution without Facial Landmarks*. <http://arxiv.org/abs/1906.07078>
- Duta, I. C., Liu, L., Zhu, F., & Shao, L. (2020). *Improved Residual Networks for Image and Video Recognition*. <http://arxiv.org/abs/2004.04989>
- Ge, S., Zhao, S., Li, C., & Li, J. (2018). *Low-resolution Face Recognition in the Wild via Selective Knowledge Distillation*. <https://doi.org/10.1109/TIP.2018.2883743>
- Gómez-Silva, M. J., Armingol, J. M., & de la Escalera, A. (2020). *Triplet Permutation Method for Deep Learning of Single-Shot Person Re-Identification*. <http://arxiv.org/abs/2003.08303>
- Grgic, M., Delac, K., & Grgic, S. (2011). SCface - Surveillance cameras face database. *Multimedia Tools and Applications*, 51(3), 863–879. <https://doi.org/10.1007/s11042-009-0417-2>
- Grm, K., Dobrišek, S., Scheirer, W. J., & Štruc, V. (2020). *Face hallucination using cascaded super-resolution and identity priors*. <http://arxiv.org/abs/1805.10938>
- Hsu, H.-J., & Chen, K.-T. (2017). DroneFace: An Open Dataset for Drone Research. *Proceedings of the 8th ACM on Multimedia Systems Conference*. <https://api.semanticscholar.org/CorpusID:26235272>
- Hu, X., Ren, W., LaMaster, J., Cao, X., Li, X., Li, Z., Menze, B., & Liu, W. (2021). *Face Super-Resolution Guided by 3D Facial Priors*. <http://arxiv.org/abs/2007.09454>
- Huang, G. B., Ramesh, M., Berg, T., & Learned-Miller, E. (2008). *Labeled Faces in the Wild: A Database for Studying Face Recognition in Unconstrained Environments*. <http://vis-www.cs.umass.edu/lfw/>
- Huang, H., He, R., Sun, Z., & Tan, T. (2019). Wavelet Domain Generative Adversarial Network for Multi-scale Face Hallucination. *International Journal of Computer Vision*, 127(6), 763–784. <https://doi.org/10.1007/s11263-019-01154-8>
- Huang, X., & Belongie, S. (2017). *Arbitrary Style Transfer in Real-time with Adaptive Instance Normalization*. <http://arxiv.org/abs/1703.06868>
- Jiang, J., Wang, C., Liu, X., & Ma, J. (2021). *Deep Learning-based Face Super-Resolution: A Survey*. <http://arxiv.org/abs/2101.03749>
- Kim, J., Li, G., Yun, I., Jung, C., & Kim, J. (2020). *Edge and Identity Preserving Network for Face Super-Resolution*. <https://doi.org/10.1016/j.neucom.2021.03.048>
- Knoche, M., Elkadeem, M., Hormann, S., & Rigoll, G. (2023). Octuplet Loss: Make

- Face Recognition Robust to Image Resolution. *2023 IEEE 17th International Conference on Automatic Face and Gesture Recognition, FG 2023*. <https://doi.org/10.1109/FG57933.2023.10042669>
- Knoche, M., Hörmann, S., & Rigoll, G. (2021). *Susceptibility to Image Resolution in Face Recognition and Trainings Strategies*. <https://doi.org/10.4230/LITES.8.1.1>
- Lai, S.-C., & Lam, K.-M. (2021). Deep Siamese network for low-resolution face recognition. *2021 Asia-Pacific Signal and Information Processing Association Annual Summit and Conference (APSIPA ASC)*, 1444–1449.
- Li, J., Bare, B., Zhou, S., Yan, B., & Li, K. (2021). Organ-Branched CNN for Robust Face Super-Resolution. *2021 IEEE International Conference on Multimedia and Expo (ICME)*, 1–6. <https://doi.org/10.1109/ICME51207.2021.9428152>
- Li, M., Zhang, Z., Yu, J., & Chen, C. W. (2021). Learning Face Image Super-Resolution Through Facial Semantic Attribute Transformation and Self-Attentive Structure Enhancement. *IEEE Transactions on Multimedia*, 23, 468–483. <https://doi.org/10.1109/TMM.2020.2984092>
- Li, P., Prieto, L., Mery, D., & Flynn, P. (2018). *On Low-Resolution Face Recognition in the Wild: Comparisons and New Techniques*. <https://doi.org/10.1109/TIFS.2018.2890812>
- Li, X., Li, W., Ren, D., Zhang, H., Wang, M., & Zuo, W. (2020). Enhanced Blind Face Restoration with Multi-Exemplar Images and Adaptive Spatial Feature Fusion. *CVPR*, 2706–2715. <https://github.com/csxmli2016/ASFFNet>.
- Li, X., Liu, M., Ye, Y., Zuo, W., Lin, L., & Yang, R. (2018). *Learning Warped Guidance for Blind Face Restoration*. <http://arxiv.org/abs/1804.04829>
- Lu, Z., Jiang, X., & Kot, A. (2018). Deep Coupled ResNet for Low-Resolution Face Recognition. *IEEE Signal Processing Letters*, 25(4), 526–530. <https://doi.org/10.1109/LSP.2018.2810121>
- Massoli, F. V., Amato, G., & Falchi, F. (2020). *Cross-Resolution Learning for Face Recognition*. <https://doi.org/10.1016/j.imavis.2020.103927>
- Mudunuri, S. P., Sanyal, S., & Biswas, S. (2018). GenLR-Net: Deep framework for very low-resolution face and object recognition with generalization to unseen categories. *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR) Workshops*, 602–611.
- Nech, A., & Kemelmacher-Shlizerman, I. (2017). *Level Playing Field for Million Scale Face Recognition*. <http://arxiv.org/abs/1705.00393>
- Sandler, M., Howard, A., Zhu, M., Zhmoginov, A., & Chen, L.-C. (2018). *MobileNetV2: Inverted Residuals and Linear Bottlenecks*. <http://arxiv.org/abs/1801.04381>
- Sapkota, A., & Boulton, T. E. (2013). Large scale unconstrained open set face database. *2013 IEEE Sixth International Conference on Biometrics: Theory, Applications and Systems (BTAS)*, 1–8. <https://doi.org/10.1109/BTAS.2013.6712756>
- Schaefer, S., Mcphail, T., & Warren, J. (2006). Image Deformation Using Moving Least Squares. *ACM SIGGRAPH*.
- Sun, J., Wenming, Y., Shen, Y., & Liao, Q. (2020). Classifier shared deep network with multi-hierarchy loss for low resolution face recognition. *Signal Processing: Image Communication*, 82, 115766. <https://doi.org/10.1016/j.image.2019.115766>
- Talreja, V., Taherkhani, F., Valenti, M. C., & Nasrabadi, N. M. (2019). *Attribute-Guided Coupled GAN for Cross-Resolution Face Recognition*. <http://arxiv.org/abs/1908.01790>
- Terhörst, P., Ihlefeld, M., Huber, M., Damer, N., Kirchbuchner, F., Raja, K., & Kuijper, A. (2021). *QMagFace: Simple and Accurate Quality-Aware Face Recognition*. <http://arxiv.org/abs/2111.13475>

- Wang, C., Jiang, J., & Liu, X. (2021). Heatmap-Aware Pyramid Face Hallucination. *2021 IEEE International Conference on Multimedia and Expo (ICME)*, 1–6. <https://doi.org/10.1109/ICME51207.2021.9428256>
- Wang, K., Oramas, J., & Tuytelaars, T. (2020). *Multiple Exemplars-based Hallucination for Face Super-resolution and Editing*. <http://arxiv.org/abs/2009.07827>
- Wolf, L., Hassner, T., & Maoz, I. (2011). Face recognition in unconstrained videos with matched background similarity. *CVPR 2011*, 529–534. <https://doi.org/10.1109/CVPR.2011.5995566>
- Yin, Y., Robinson, J. P., Zhang, Y., & Fu, Y. (2020). *Joint Super-Resolution and Alignment of Tiny Faces*. <http://arxiv.org/abs/1911.08566>
- Yu, X., Zhang, L., & Xie, W. (2021). Semantic-Driven Face Hallucination Based on Residual Network. *IEEE Transactions on Biometrics, Behavior, and Identity Science*, 3(2), 214–228. <https://doi.org/10.1109/TBIOM.2021.3051268>
- Zangeneh, E., Rahmati, M., & Mohsenzadeh, Y. (2020). *Low Resolution Face Recognition Using a Two-Branch Deep Convolutional Neural Network Architecture*. <http://arxiv.org/abs/1706.06247>
- Zeng, D., Chen, H., & Zhao, Q. (2016). Towards resolution invariant face recognition in uncontrolled scenarios. *2016 International Conference on Biometrics (ICB)*, 1–8. <https://doi.org/10.1109/ICB.2016.7550087>
- Zha, J., & Chao, H. (2019). TCN: Transferable Coupled Network for Cross-Resolution Face Recognition. *ICASSP 2019 - 2019 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, 3302–3306. <https://doi.org/10.1109/ICASSP.2019.8682384>
- Zhao, X. (2021). *Homogeneous Low-Resolution Face Recognition Method based Correlation Features*. <http://arxiv.org/abs/2111.13175>