



Application of GAN Deblurring Algorithm in Face Recognition

Nan Deng^{1,2*}, Ling Pu³, Liye Tao¹, Lingli Guo⁴, Runtao Yang⁵, Jianfang Lin⁶

¹ School of Mathematics and Computer Science, Hebei Minzu Normal University, Chengde 067000, China

² The Technology Innovation Center of Cultural Tourism Big Data of Hebei Province, Chengde 067000, China

³ Bureau of Public Works of Shenzhen Municipality Engineering Management Center, Shenzhen 518000, China

⁴ Department of mathematics, Changzhi University, Changzhi 046000, China

⁵ Department of Artificial Intelligence, Tangshan University, Tangshan 063000, China

⁶ Department of Computer Science and Technology, Tangshan Normal University, Tangshan 063000, China

Corresponding Author Email: dengnan@hbun.edu.cn

Copyright: ©2025 The authors. This article is published by IETA and is licensed under the CC BY 4.0 license (<http://creativecommons.org/licenses/by/4.0/>).

<https://doi.org/10.18280/ts.420519>

ABSTRACT

Received: 2 February 2025

Revised: 19 July 2025

Accepted: 23 September 2025

Available online: 31 October 2025

Keywords:

facial recognition, generative adversarial network, convolutional neural network, dual-discriminator architecture, wiener filtering

With the advancement of informatization, many identity authentication processes have begun to be conducted online, thus putting forward new requirements for the performance of face recognition models. Most existing image deblurring algorithms based on deep neural networks have problems such as a large number of parameters and an insignificant improvement in face recognition accuracy. For these reasons, this research designs a face recognition algorithm based on the generative adversarial network (GAN) framework and deblurring principles. The model uses GAN as the core algorithm, with improvements and optimizations made to both the generator module and the discriminator module. A preprocessing layer is added to the Convolutional Neural Network (CNN) of the generator module. In the discriminator module, a dual-discriminator architecture is proposed, which includes global discriminator and local discriminator algorithms. These improvement and optimization methods aim to address issues such as the insignificant improvement in accuracy in scenarios related to blurred face recognition. Meanwhile, to tackle the model's poor capability in processing blurred images, the research introduces filtering technology to perform deblurring operations on images. The algorithm model first performs filtering processing on images and then uses the improved GAN algorithm for feature extraction, which can further enhance the model's ability to recognize blurred images. Through testing, the model achieves an average recall rate of 88.12% on the LFW dataset; the accuracy rates on the LFW and CASIA WebFace datasets reach 81.32% and 79.81% respectively; and the F1-scores reach 95.27% and 94.73% respectively. The experimental results show that the proposed model utilizes more overall and detailed features of human faces to address the problem of blurred face recognition, thereby improving the recognition accuracy.

1. INTRODUCTION

The advent of the information age has brought potential risks to information security, making information recognition technology particularly important nowadays [1]. As the most commonly used identity information recognition method, face recognition has been applied in mobile payment, smart door locks, identity verification, and other fields [2, 3]. Conventional face recognition technology has developed quite maturely; for example, manufacturers like SenseTime hold numerous high-end technical patents for recognition tasks with different precision requirements. Clear and high-quality facial image data can help algorithms perform their functions better, while blurred and low-quality facial image data will have a significant adverse impact on the output results of algorithms [4, 5]. Generative Adversarial Network (GAN), as an emerging deep learning model, has achieved remarkable results in tasks such as image super-resolution reconstruction and image recognition. In this study, a face recognition

algorithm is designed based on the generative adversarial network framework and deblurring principles.

The innovations of this research are as follows: The first is the improvement of the generation module. The image preprocessing layer is integrated into the CNN generation module of the GAN algorithm. It extracts the basic features of blurred faces through convolution, then extracts CNN convolutional features, and finally fuses the basic features with the convolutional features. This makes full use of the CNN algorithm to improve image quality while retaining valuable information such as the position and type of blurred regions. The second is the improvement of the discrimination module. A dual-discriminator architecture is proposed, which includes a global discriminator and a local discriminator. The global discriminator decomposes large convolution kernels into 6 small-scale convolutions, adds a cross-scale feature fusion layer, and adjusts the channels through 1×1 convolution before summing them up. This avoids feature fragmentation and improves the integrity of global features. The local

discriminator introduces a key region priority strategy, dynamically adjusting the penalty loss weights of sensitive regions such as eyes and nose, thus simplifying parameter tuning. The global discriminator strengthens the texture and structural information of the entire image, while the local discriminator enhances key facial information such as eyes, nose, and mouth. Through their synergistic effect, the ability to distinguish the authenticity of generated images is significantly improved. The third is the construction of a blurred face recognition method that integrates Wiener filtering-based image deblurring with the above-improved GAN algorithm. First, Wiener filtering is applied to blurred images, and then the filtered images are fed into the GAN network to complete face recognition.

The article is divided into four parts. The first part covers related work such as literature review. The second part is the research method, which mainly focuses on the study of face recognition models based on adversarial theory and image deblurring processing methods. The third part is performance verification, where the model's performance is verified through multiple sets of experiments. The fourth part is the conclusion, which summarizes and analyzes the relevant experimental data.

2. RELATED WORKS

The key to online identity verification lies in achieving efficient and secure face recognition. In the face of various complex scenarios, researchers have been continuously exploring and improving face recognition technologies. Lu et al. [6] introduced a refined CNN framework trained on enriched datasets, targeting performance gaps observed in conventional facial recognition datasets. This model transformed faces multiple times and recombined their features to expand the dataset. When compared with other face recognition models, this model showed good stability and progressiveness. Montero et al. [7] analyzed and studied the face recognition problem of wearing masks in the context of the COVID-19 global pandemic, and proposed an end-to-end face recognition model training model based on the ArcFace architecture. It included a modification module for the underlying architecture and loss function calculation, as well as a dataset expansion module. These experiments confirmed that the recognition performance of this model for detecting facial images with masks was superior to the baseline model, and it achieved an average accuracy of 99.78% on the original dataset. Xie et al. [8] proposed a GAN-based image fusion algorithm to preserve texture and target information in the source image. It used adaptive gradient decomposition algorithm to extract image features, separates high-frequency and low-frequency components, and then used principal component analysis for fusion. The proposed model was trained on the TNO and RoadScene datasets and showed good stability and high recognition performance. Fu et al. [9] proposed a Dual Variational Generation (DVG-Face) framework based on the HFR formula to address the shortcomings of differentiation and insufficient heterogeneous data in the heterogeneous facial recognition. This framework was based on a dual variational generator to learn the joint distribution of paired heterogeneous images, and incorporated identity preservation parameters to ensure identity consistency in heterogeneous images. Finally, this architecture achieved better performance than the most advanced methods on seven

challenging databases belonging to five HFR tasks, reflecting its progressiveness nature. Research led by Xiao et al. [10] revealed significant susceptibility of deep neural architectures to physically realizable patch-based attacks. In response, a novel defense paradigm was developed through adversarial feature engineering on dimensionally reduced data manifolds. By leveraging facial embeddings as perturbation sources, the regularized model exhibited enhanced attack transferability across heterogeneous platforms during empirical validation, coupled with measurable improvements in recognition performance benchmarks. To address data reliability issues in facial recognition systems, Shang et al. [11] proposed a unified uncertainty-aware architecture that dynamically adjusts sample learning weights. This framework achieves state-of-the-art recognition accuracy with reduced computational overhead in benchmark evaluations.

The widespread use of masks during the COVID-19 pandemic has posed significant challenges to facial biometric systems. In response, Hariri [12] developed a method combining occlusion removal with deep learning-based feature extraction. This approach removes occluded areas, focuses on extracting features from the eyes and forehead using deep neural networks, and uses a fully connected layer and multi-layer perceptrons for classification. Experiments show this method improves recognition accuracy and reliability compared to other advanced techniques.

Qiu et al. [13] designed an occlusion-robust visual perception framework using end-to-end learned representations. The architecture decomposes facial saliency maps with cascaded residual transformers and employs learnable deformation fields for context-aware feature reconstruction. On the Megaface Challenge 1 dataset, it achieves a 4.7% improvement in Rank-1 accuracy under 60% occlusion compared to state-of-the-art baselines.

Zhang et al. [14] found that different face regions affect recognition performance. They proposed a visual attention orchestration framework using reinforcement learning policies to control deep convolutional pathways. This model integrates attention mechanisms with feature networks and applies regularization to improve efficiency. It has shown promising results on major face verification databases, confirming its feasibility.

Terh rst et al. [15] highlighted the significant impact of facial recognition on key decisions. To enhance trustworthiness, they analyzed the effects of 47 attributes on two mainstream facial recognition systems. Based on their findings, they introduced a new method that effectively reduces bias and improves recognition performance.

To sum up, research on face recognition based on various machine learning methods has achieved certain results at this stage. Adversarial learning is mainly used for image reconstruction, and there are relatively few studies applying it to face recognition. Moreover, traditional GANs still fail to meet the requirements of face recognition in terms of recognition speed. Therefore, this research aims to improve GANs and construct an efficient face recognition model based on the improved GAN.

3. CONSTRUCTION OF A FACIAL RECOGNITION MODEL BASED ON IMPROVED GAN AND DEBLURRING PROCESSING

The key to various information authentication scenarios lies

in achieving efficient and secure face authentication. In this study, an efficient face recognition model is constructed based on the GAN algorithm and deblurring processing technology.

3.1 Face recognition based on the enhanced GAN algorithm

GAN technology (Generative Adversarial Networks) is a deep learning-based generative model proposed by Canadian scientist Ian Goodfellow and his team in 2014. Through the "adversarial training" process of two neural networks, it learns the distribution of real data, thereby generating new samples similar to real data [16]. The generation module first accepts fuzzy images and random noise as inputs. Fuzzy face images provide background information, while random noise is used for the randomness of the generator, making the generated images diversified. In the module, multiple convolution layers and activation functions are used to extract the features of fuzzy images and feature pyramid technology is used to fuse the feature information of different levels to ensure that the generated images have sufficient details both globally and locally. Finally, through a set of up-sampling and convolution operations, the extracted features are reconstructed into clear face images. The discriminative pathway executes instance-level differentiation via cascaded convolutional operations with residual gating mechanisms [17]. The input of the discriminant module is the real face image and the generated face image. The features of input images are extracted by multiple convolution layers, and a series of nonlinear transformations are carried out to enhance the capability of feature extraction. Finally, the module uses a fully connected layer to output the judgment result and determine whether the input is a real image. GAN-based face recognition algorithms leverage "generative" and "adversarial" traits to boost performance. The generator takes random noise or low-dimensional vectors, using multi-layer neural networks to produce "fake data" resembling real faces, aiming to "deceive" the discriminator [16]. The discriminator receives real or generated faces, outputs a 0-1 probability to judge authenticity, seeking accurate distinction. Generator and discriminator compete continuously: the generator creates more realistic faces to fool the discriminator, which hones its distinguishing ability. Eventually, a dynamic equilibrium emerges — generated data closely matches real data distribution, with the discriminator's accuracy near 50%. In this process, the generator produces numerous realistic face samples (varied poses, expressions, lighting, occlusions), enriching training sets to prevent overfitting. The discriminator learns key features like facial structure when distinguishing real from fake; the generator grasps essential facial traits (e.g., feature positional constraints) to ensure realism. Thus, GAN captures more abstract, representative facial features (not superficial noise like lighting), enhancing feature distinguishability. Training alternates: first, fix the generator, input real and virtual faces to the discriminator, which uses labels (1 for real, 0 for generated) to calculate losses and update parameters for better distinction. Then fix the discriminator; the generator inputs virtual faces, striving for a "real" judgment (label 1), updating parameters to boost realism. This cycle continues until equilibrium [17]. Figure 1 delineates the adversarial architecture comprising.

The formulation of the loss function critically determines the convergence properties and synthesis quality in generative adversarial networks, with the fundamental objective

expressed as Eq. (1).

$$M(G_{en}, D_{is}) = \sum_{x_{re}}^{p_{re}} [\log D_{is}(x_{re})] + \sum_{x_{re}}^{p_{di}} [\log(1 - D_{is}(x_{re}))] \quad (1)$$

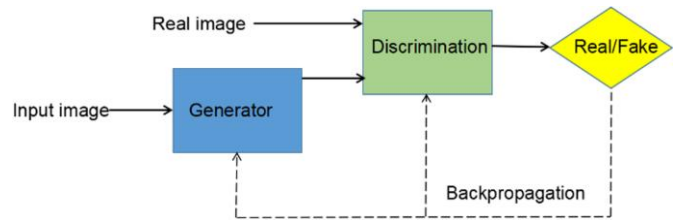


Figure 1. GAN model architecture

where, x_{re} refers to the input real sample data. p_{re} represents the true distribution of data. $M(G_{en}, D_{is})$ refers to the global loss function. G_{en} is the module of generation. p_{di} denotes the generated data distribution. D_{is} denotes the module of discrimination. $D_{is}(x_{re})$ means the output result of the discrimination module. The loss function of the generator module, after optimization, is denoted by Eq. (2). Attaining the optimal parameters specified in Eq. (1) across both the generation and discrimination modules concurrently presents significant difficulties. As a result, decomposing the function and optimizing its generation and discrimination modules on their own becomes essential. The optimized form of the loss function specific to the generator module is denoted by Eq. (2).

$$\text{Min}_{ge} M(G_{en}, D_{is}) = \sum_{x_{re}}^{p_{di}} [\log(1 - D_{is}(x_{re}))] \quad (2)$$

where, Min_{ge} represents the minimum value of the loss function in the generation module. Within the generation module, there exists an inverse proportionality between this value and the effectiveness of sample generation. Regarding the recognition module, its optimized loss function is denoted by the Eq. (3).

$$\begin{aligned} \text{Max}_{di} M(G_{en}, D_{is}) &= \sum_{x_{re}}^{p_{re}} [\log D_{is}(x_{re})] \\ &+ \sum_{x_{re}}^{p_{di}} [\log(1 - D_{is}(x_{re}))] \end{aligned} \quad (3)$$

The design inspiration of the GAN algorithm stems from the two-player zero-sum game in game theory, and Nash equilibrium, a key concept in game theory, describes a stable state where the strategies of all participants in the game reach a balance [18]. In this state, no participant can improve their own payoff by unilaterally changing their strategy. In GAN, the generator and the discriminator are like the two sides of a game: the generator strives to generate fake data that is difficult for the discriminator to distinguish, while the discriminator tries its best to accurately judge the authenticity of the input data. Their adversarial process can be seen as a search for a state similar to Nash equilibrium. However, to achieve the optimal training level, further optimization of the network is required [19].

CNN is a deep learning model widely used in image processing and computer vision fields. Thanks to its

advantages of local receptive fields, weight sharing, and hierarchical feature extraction, CNN can effectively capture spatial structure information in input data while reducing computational complexity. Using CNN for super-resolution reconstruction of facial images is currently the most widely applied method. In this study, the generator adopts a CNN network, and the CNN convolution kernels of the generator are improved using an image preprocessing layer [20]. The input of the generator module is a relatively clear facial image after deblurring processing. The design of the CNN network is as follows: first, 3 convolutional layers (CONV), batch normalization (BN), and ReLU activation function are used to extract low-level and mid-level features. Second, it enters a residual network (ResNet) containing 5 residual blocks, each consisting of 2 convolutional layers. Through skip connections, details are retained and gradient vanishing is prevented, thus maintaining good performance and stability as the depth increases [5]; finally, features are reconstructed through 3 deconvolutional layers, and the generated image is output. The CNN network mainly processes the generated clear faces, but the original blurred faces still contain valuable information (such as the position of blurred regions and features of blur types). Therefore, in this study, first, basic features of the original blurred face are extracted through simple convolution, then image convolutional features are extracted, and finally, the basic features are fused with features from different convolutional layers.

The network adopts an encoder-decoder architecture and designs a complete process of feature extraction, fusion, and reconstruction for the task of blurred face clarification, with specific implementation divided into two parts: encoder and decoder.

The encoder first extracts key low-frequency information through a preprocessing layer: the first layer uses 3×3 convolution (64 channels) combined with ReLU to capture basic structural features such as edges and contours of blurred faces; the second layer adjusts the number of channels through 1×1 convolution to match subsequent modules. Then it enters the main feature extraction network: first, 3 convolutional layers (each containing convolution, BN, and ReLU) are used to extract low-level to mid-level features, and then 5 residual

blocks (each containing 2 convolutional layers) are used for deep feature mining. The skip connections of the residual blocks not only retain detailed information but also alleviate the problem of gradient vanishing. During the encoding process, the low-frequency features of the preprocessing layer and the output features of convolutional layers at various stages are fused through channel concatenation, ensuring that basic structural information is not excessively modified.

The decoder adopts a 3-layer deconvolution structure symmetrical to the encoder to realize stepwise reconstruction of features: each deconvolution layer first performs upsampling through bilinear interpolation combined with 3×3 convolution (to avoid checkerboard effects), gradually restoring the feature map size to the input image size; at the same time, each layer is channel-concatenated with the fused features of the corresponding level in the encoder to form cross-level feature complementation (e.g., the shallow layer of the decoder fuses high-level semantic features from the deep layer of the encoder, while the deep layer of the decoder focuses on fusing low-frequency features from the preprocessing layer). After concatenation, 1×1 convolution is used to compress channels and reduce redundancy, followed by 3×3 convolution to refine features. Finally, the number of channels is adjusted to 3 (RGB channels) through the last convolution layer, and the Sigmoid activation function is used to output the clarified face image, achieving the dual goals of "retaining original structure + supplementing high-frequency details". The structure of the improved generation module is shown in Figure 2.

The computational time of the discriminant module accounts for a large proportion of the total operation time of the model, so it is necessary to optimize it to improve the overall running speed of the model [21]. We propose a dual-discriminator architecture, which includes the design of a global discriminator and a local discriminator. The global discriminator strengthens the texture and structural information of the entire image, while the local discriminator enhances key facial information such as eyes, nose, and mouth. Through their synergistic effect, the ability to distinguish the authenticity of generated images is significantly improved.

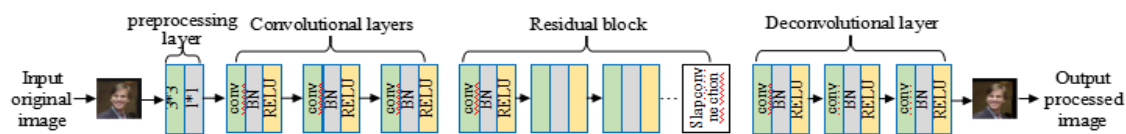


Figure 2. Improved generator structure

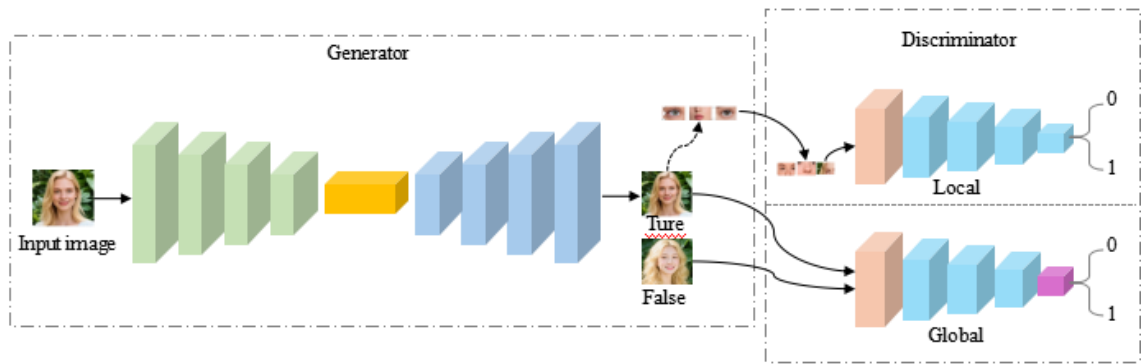


Figure 3. A-GAN model structure

The design of the global discriminator is achieved by decomposing a large convolution kernel into six small-scale

convolution kernels and adding one cross-scale feature fusion layer among the six decomposed small-scale convolutional layers. Specifically, the output features from different small convolutions are adjusted in terms of channels through 1×1 convolution and then summed. This avoids feature fragmentation caused by the independent operation of small convolutions, and while maintaining the advantage of computational efficiency, it enhances the integrity of global features.

The design of the local discriminator involves introducing a key region priority strategy during the random segmentation of facial images. It prioritizes segmenting sensitive regions for recognition, such as the eyes and nose, and dynamically adjusts the weights of local penalty losses according to the importance of regions—assigning higher weights to the eye and nose regions. These adjustments do not require manually setting fixed weight groups; instead, parameter tuning can be completed by simplifying parameters. The overall structure diagram of the improved generation module and discrimination module is shown in Figure 3.

At this point, the GAN model based on the improved mechanism is constructed, referred to as A-GAN for short. Its advantage lies in the discriminator's ability to enhance the joint discrimination capability for global structures and local details through feature fusion and dynamic weights.

3.2 Development of an adversarial learning-based face recognition model incorporating deblurring

To address the recognition of blurred face images, deblurring technology is introduced to further improve the model. The principles of image deblurring are generally similar [22]. In this study, the Wiener filtering algorithm is adopted for deblurring processing of face images.

Wiener filtering, also known as minimum mean square error filtering, operates on the principle that image restoration is achieved by solving for the estimated value $\hat{f}(x, y)$ when the mean square error between the estimated value $\hat{f}(x, y)$ of the original clear image and the actual image $f(x, y)$ is minimized. The mean square error between $f(x, y)$ and its estimated value $\hat{f}(x, y)$ is:

$$e^2 = E\{(f(x, y) - \hat{f}(x, y))^2\} \quad (4)$$

where, e^2 represents the mean square error value, and there is also:

$$f(x, y) = F^{-1}[F(u, v)] = F^{-1}\left[\frac{G(u, v) - N(u, v)}{H(u, v)}\right] \quad (5)$$

Therefore, when the value of Eq. (5) is minimized, it represents the value of image restoration via Wiener filtering. After Fourier transform, $\hat{f}(x, y)$ is expressed as:

$$\hat{F}(u, v) = \left[\frac{1}{H(u, v)} \frac{|H(u, v)|^2}{|H(u, v)|^2 + \frac{P_n(u, v)}{P_f(u, v)}} \right] G(u, v) \quad (6)$$

where, $G(u, v)$ is the original clear image. $g(x, y)$ is the result after Fourier transform. $P_f(u, v)$ is the power spectrum of the original image. $P_n(u, v)$ is the power spectrum of the noise. $P_n(u, v)/P_f(u, v)$ is the signal-to-noise ratio.

Let $K = P_n(u, v)/P_f(u, v)$, then Eq. (6) can be transformed into:

$$\hat{F}(u, v) = \left[\frac{1}{H(u, v)} \frac{|H(u, v)|^2}{|H(u, v)|^2 + K} \right] G(u, v) \quad (7)$$

where, $H(u, v)$ is the point spread function. K is the signal-to-noise ratio. As can be seen from Eq. (7), both $H(u, v)$ and K are of great importance in the image restoration process using Wiener filtering. When a blurred image is free from noise interference, Wiener filtering yields good results in image restoration. However, for noisy blurred images, if the value of K is uncertain (for example, $K = 0$), Wiener filtering is basically unable to complete the image restoration. Conversely, if an accurate value of K is obtained, Wiener filtering can accomplish the image restoration task effectively. In this regard, the values of K and $H(u, v)$ restrict the application scenarios of Wiener filtering. The Wiener filtering restoration algorithm features simple computation and a clear structure, but it has drawbacks: it cannot automatically identify blur parameters and has poor dynamic performance. The effect of Wiener filtering is shown in Figure 4.

The overall process of deblurring a blurred image using Wiener filtering is as follows: obtain the blurred image to be processed, determine the point spread function (PSF), convert the image and the PSF to the frequency domain, and use the Fourier transform to convert the convolution operation in the spatial domain into a multiplication operation in the frequency domain to simplify the computation. Then, estimate the noise power spectrum and the signal power spectrum, apply the Wiener filtering formula to calculate the frequency-domain result of the restored image, convert the frequency-domain result back to the spatial domain, and optionally fine-tune the restored result.

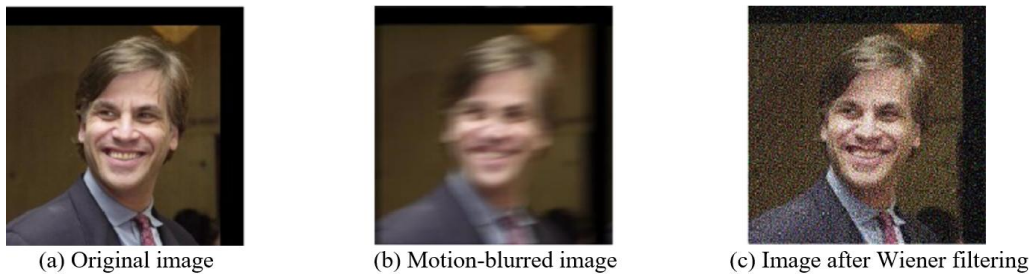


Figure 4. Wiener filtering effect diagram

In summary, the core of the Wiener filtering process is to combine the PSF and noise characteristics in the frequency domain, suppress blurring and reduce noise interference

through the filtering formula, and finally obtain the deblurred image through inverse transformation. Its effectiveness highly depends on the accuracy of the Point Spread Function and the

estimation precision of the noise power spectrum.

The algorithm first applies Wiener filtering to the blurred face image, then resizes the filtered image to meet the input size requirements of the GAN model. Next, the GAN model generates a high-resolution image, which serves as the restored face image. Finally, face recognition is performed on this restored image. The overall algorithm, integrating blurred

image processing and the improved GAN, is thus constructed and named AD-GAN. The overall flowchart of the AD-GAN model is shown in Figure 5. After face deblurring processing, the model has been optimized, leading to improvements in both the algorithm's processing efficiency and the model's recognition accuracy.

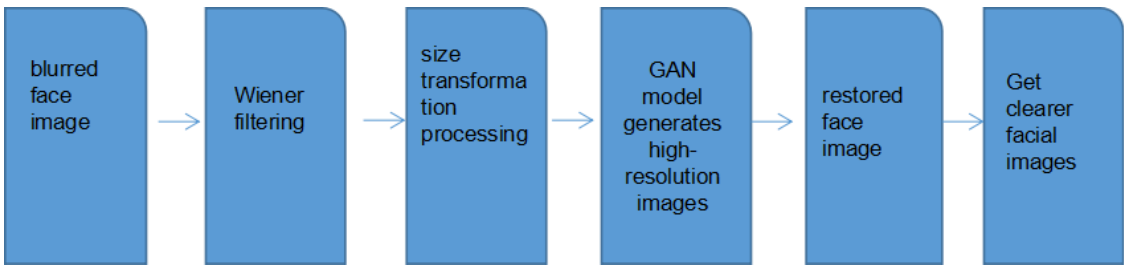


Figure 5. Overall flowchart of the AD-GAN model

4. PERFORMANCE VALIDATION OF THE ADVERSARIAL LEARNING FACE RECOGNITION MODEL INTEGRATED WITH DEBLURRING PROCESSING

The training and testing in this experiment use datasets including LFW and CASIA WebFace. The CASIA WebFace face dataset, released by the Chinese Academy of Sciences, is a large-scale face dataset currently widely used in tasks such as face recognition and identity verification. To verify the model's performance in regional framing, segmentation, and recognition, this study randomly selected a face image as the input to the AD-GAN model, and the output results of the model are shown in Figure 6. It can be seen from Figure 6 that the AD-GAN model performs well in face localization and segmentation, and Figure 6(c) can also accurately locate the key points of the face.

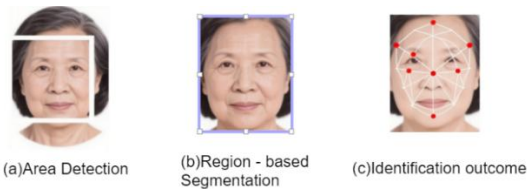


Figure 6. The algorithmic performance of the AD-GAN model

Face recognition speed is crucial. To test the recognition speed of the AD - GAN model, in this research, A - GAN and GAN are used as controls. 100 images each with pixel sizes of 720*720, 360*360, and 180*180 are selected from CASIA WebFace and IMDB-WIKI as inputs. Figure 7 documents the output times of the three models. In Figure 7(a), the image processing time of AD-GAN was shorter than that of A-GAN and GAN, while there was not much difference in the output time between the A-GAN and GAN models. In Figure 7(b), A-GAN and GAN models' output time had a significant difference, but the output time of AD-GAN was still the lowest. The average processing time of AD-GAN for 180*180, 360*360, and 720*720 images was 46.23s, 59.84s, and 71.37s, respectively.

AD-GAN has added a deblurring processing module, which also has good recognition accuracy for blurred images. To

confirm the deblurring performance of AD-GAN, images from the LFW dataset with blur parameters of 0.8 and 0.4 were chosen as inputs, as shown in Figure 8. As indicated in Figure 8(a), the average blurriness of the output images generated by AD-GAN, A-GAN, and GAN dropped to 0.171, 0.226, and 0.360 respectively. From Figure 8(b), the average blurriness of the output images produced by AD-GAN, A-GAN, and GAN decreased to 0.082, 0.045, and 0.051. It is evident that AD-GAN exhibits a more remarkable deblurring effect when the blurriness is high. When the fuzziness is low, the processing effect of these three models is not significantly different.

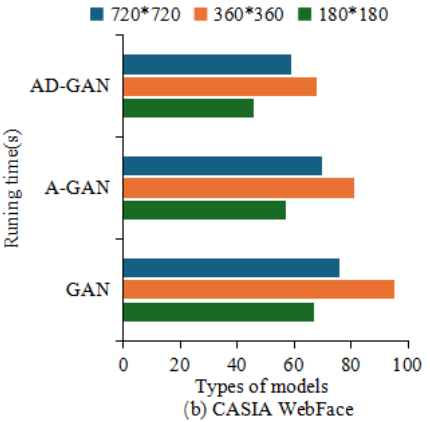
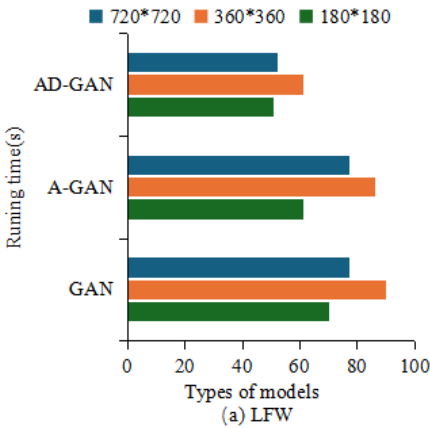


Figure 7. Display diagram of recognition time for different models

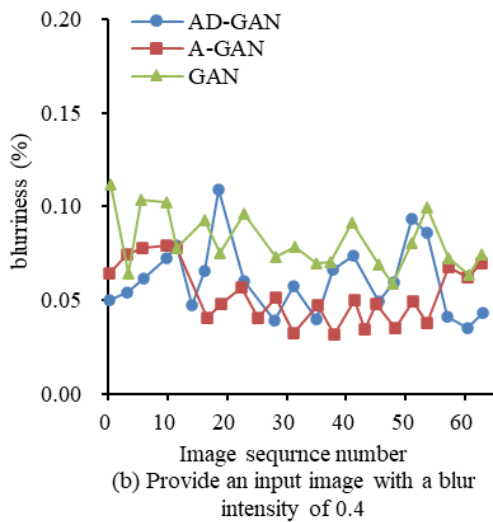
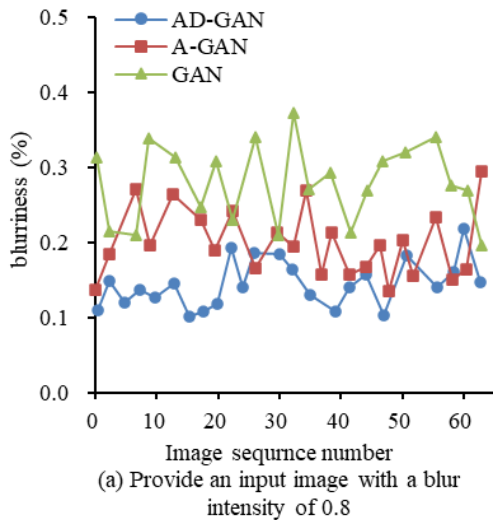


Figure 8. Presentation of deblurring performances across various models

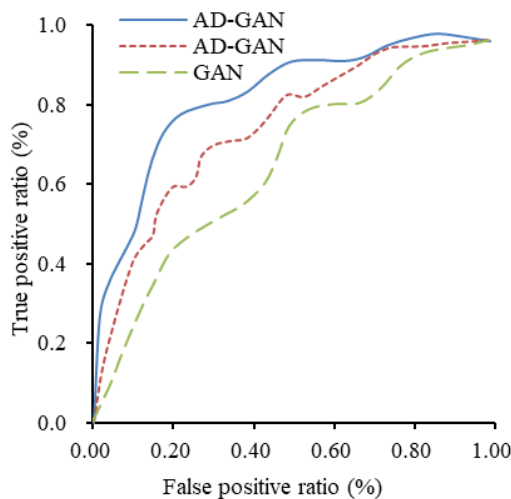


Figure 9. Schematic diagram of ROC curves for deblurring performance demonstration of diverse models using the LFW dataset

ROC is a variation curve obtained under specific conditions for different stimuli. Using the LFW dataset as input for AD-GAN, A-GAN, and GAN, Figure 9 shows the ROC of the three models. When the true positive rate of this model sample

approached 100%, the false positive rate also approached 1. The true positive rate of AD-GAN approached 100% faster, indicating that the model was most sensitive to stimuli, followed by A-GAN and GAN.

Precision stands as the key metric for a model, offering a straightforward reflection of its capability in facial recognition. For this research, the CASIA WebFace dataset and LFW dataset served as input sources, with the focus on documenting how model training duration correlates with precision. The findings of the experiment are illustrated in Figure 10. From Figure 10(a), it is evident that the AD-GAN model achieves rapid convergence on the CASIA WebFace dataset, and post-convergence, its precision ranks the highest among the three models. Looking at Figure 10(b), on the LFW dataset, the AD-GAN model requires a marginally longer time to converge compared to the A-GAN model; however, its precision still remains the top among the three. After converging on the CASIA WebFace and LFW datasets, the AD-GAN model attains precision rates of 81.32% and 79.81% respectively.

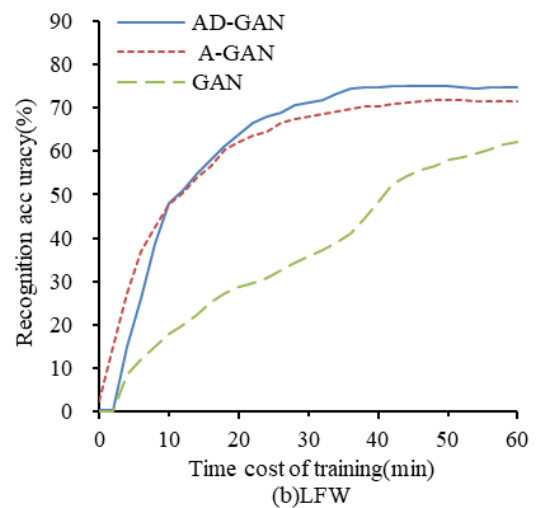
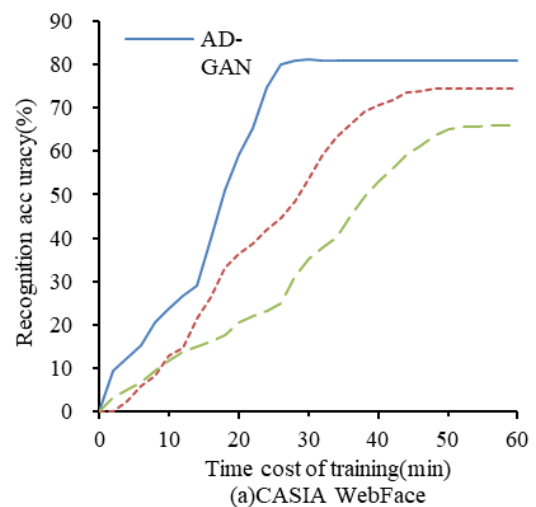


Figure 10. Relationship between training time and accuracy

Recall rate refers to the ratio of correctly predicted samples by the model to the total number of samples, and a high recall rate indicates a high probability of successful model recognition. This recall test used the LFW dataset as input to compare the recall rates of AD-GAN, A-GAN, and GAN in Figure 11(a). The above steps were repeated three times to

obtain the average recall rates in Figure 11(b). In summary, on LFW, AD-GAN had the highest average recall rate, reaching 88.12%.

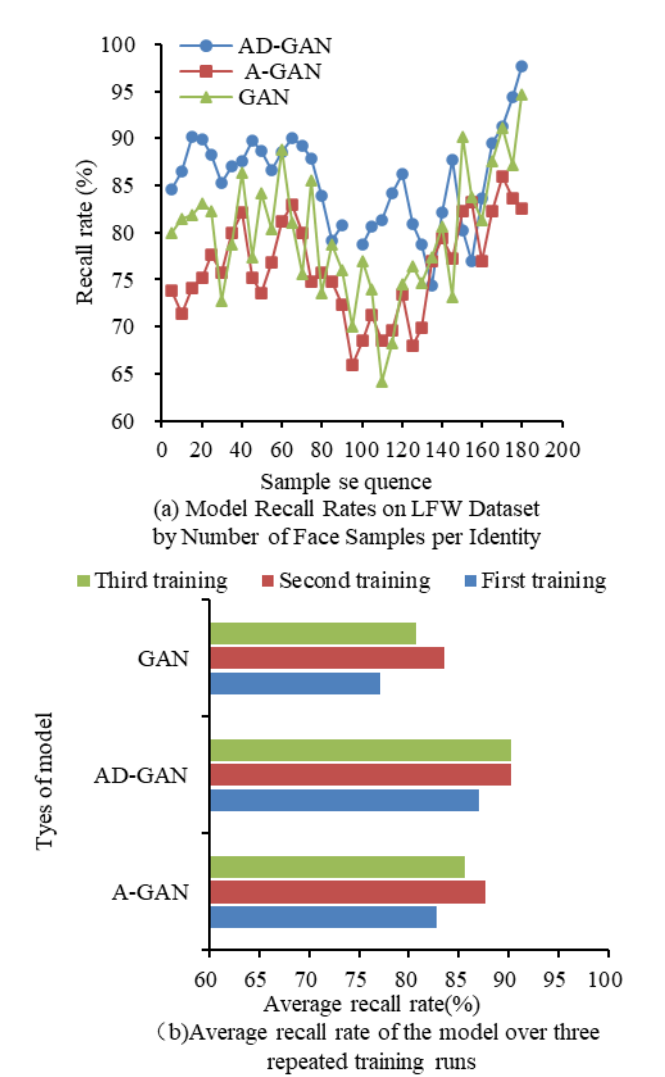


Figure 11. Graph comparing recall rates across diverse models

F1 is a comprehensive evaluation index based on the accuracy and recall of the model. Generally, if the model F1 is high, its overall performance is good. The experiment used CASIA WebFace and LFW as inputs for AD-GAN, A-GAN, and GAN. Figure 12 calculated the F1 value and training time of each model. Figure 1 (a) denotes the variation in F1 values of the three models on CASIA WebFace as training proceeds. When the training duration was 50 minutes, the F1 of AD-GAN, A-GAN, and GAN were 94.73%, 92.79%, and 87.60%, respectively. Figure 12(b) represents the variation in F1 values of the three models on IMDB-WIKI as training proceeds. When the training duration was 50 minutes, the F1 of AD-GAN, A-GAN, and GAN were 95.27%, 91.06%, and 84.91%, respectively. Overall, AD-GAN achieved higher F1-scores than the two control models, which indicates its superior overall performance.

Based on the experimental results, it is concluded that the model proposed in this study has relatively high detection accuracy and relatively stable F1-scores, and can also achieve relatively ideal results in processing blurred images. On the LFW dataset, the model with data processed by deblurring can significantly enhance the recognition rate of blurred images.

Moreover, the precision of images featuring blurriness sees a marked enhancement following processing.

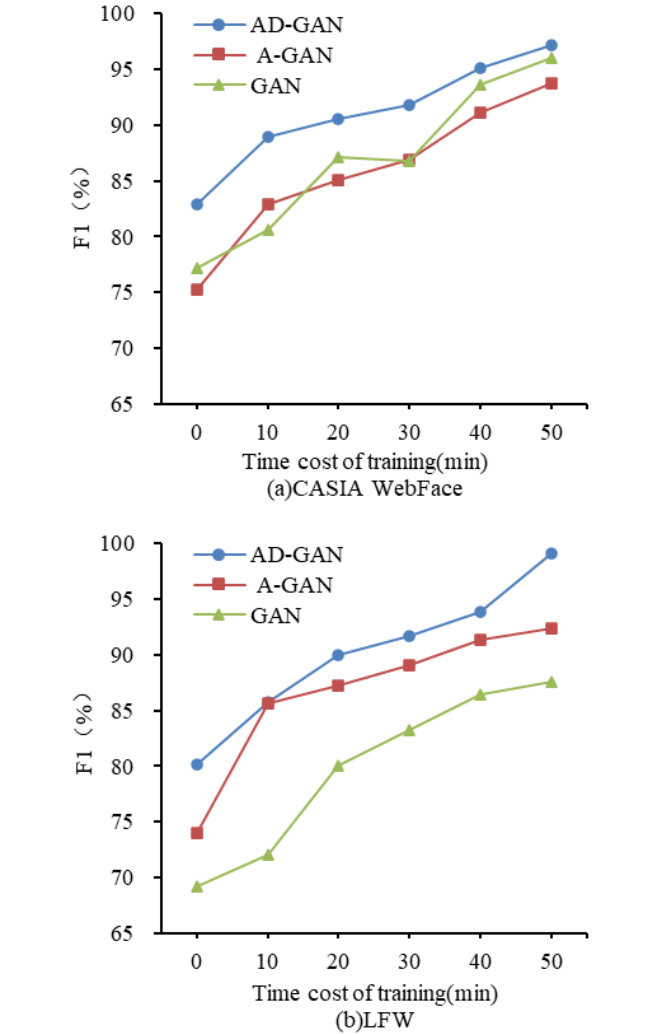


Figure 12. Diagram that compares F1-scores across diverse models

5. CONCLUSION

Given that GAN models struggle with inefficient recognition of blurred images, this research puts forward an enhanced GAN-based face recognition model that incorporates an image deblurring algorithm. This model not only retains the efficient recognition capability of GAN but also enhances the recognition performance for blurred faces by leveraging deblurring technology. The model's performance was verified on the LFW and CASIA WebFace datasets. The results of experiments denote that the model of AD-GAN gets an average recall rate of 88.12% on the LFW dataset; after convergence, the accuracy rates on the LFW and CASIA WebFace datasets reach 81.32% and 79.81% respectively, with F1-scores of 95.27% and 94.73% respectively. Meanwhile, for images with a blur degree of 0.8, the AD-GAN model reduces the blur degree to only 0.171 after deblurring, a decrease of 0.629; When dealing with input images that have a blur intensity of 0.4, the resulting output images exhibit a blur degree reduction of 0.318. The proposed model also denotes faster processing capabilities for images of different sizes compared to the comparison model: The

average processing duration of the AD-GAN model for 180×180 pixel images is 46.23 seconds; for 360×360 pixel images, it is 59.84 seconds; and for 720×720 pixel images, it is 71.37 seconds. Compared with the A-GAN model, which has average processing times of 52.44 seconds, 65.89 seconds, and 77.64 seconds respectively, AD-GAN shows obvious advantages. Future research can introduce GAN networks with stronger discrimination capabilities and faster processing speeds for further improvement to achieve better face recognition results. The findings derived from this model hold substantial importance in boosting both the precision and resilience of face recognition systems, offering fresh perspectives and approaches to guide subsequent research endeavors.

FUNDINGS

This research was supported by the Special Science and Technology Plan Project for Applied Technology Research and Development and Sustainable Development Goals Innovation Demonstration Zone of Chengde City, Hebei Province (Grant No.: 202404B025), the R&D Investment Intensity Growth Reward Fund from the Science and Technology Bureau of Chengde City (Grant No.: 202304B071), and Hebei Minzu Normal University (Grant No.: DR2023003).

REFERENCES

- [1] Zhou, K., Tang, J. (2021). Harnessing fuzzy neural network for gear fault diagnosis with limited data labels. *The International Journal of Advanced Manufacturing Technology*, 115(4): 1005-1019. <https://doi.org/10.1007/s00170-021-07253-6>
- [2] Sundaram, S.M., Narayanan, R. (2023). Human face and facial expression recognition using deep learning and SNet architecture integrated with BottleNeck attention module. *Traitement du Signal*, 40(2): 647-655. <https://doi.org/10.18280/ts.400223>
- [3] Ali, N.S., Alsafo, A.F., Ali, H.D., Taha, M.S. (2024). An effective face detection and recognition model based on improved YOLO v3 and VGG 16 networks. *International Journal of Computational Methods and Experimental Measurements*, 12(2): 107-119. <https://doi.org/10.18280/ijcmem.120201>
- [4] Hasanvand, M., Nooshyar, M., Moharamkhani, E., Selyari, A. (2023). Machine learning methodology for identifying vehicles using image processing. *Artificial Intelligence and Applications*, 1(3): 154-162. <https://doi.org/10.47852/bonviewAIA3202833>
- [5] Diba, M., Khosravi, H. (2024). SNResNet: A new architecture based on SqNxt blocks and rish activation for efficient face recognition. *Traitement du Signal*, 41(2): 949-959. <https://doi.org/10.18280/ts.410235>
- [6] Lu, P., Song, B., Xu, L. (2021). Human face recognition based on convolutional neural network and augmented dataset. *Systems Science & Control Engineering*, 9(sup2): 29-37. <https://doi.org/10.1080/21642583.2020.1836526>
- [7] Montero, D., Nieto, M., Leskovsky, P., Aginako, N. (2022). Boosting masked face recognition with multi-task arface. In 2022 16th International Conference on

- Signal-Image Technology & Internet-Based Systems (SITIS), Dijon, France, pp. 184-189. <https://doi.org/10.1109/SITIS57111.2022.00042>
- [8] Xie, Y., Liu, G., Xu, R., Bavirisetti, D.P., Tang, H., Xing, M. (2023). R2F-UGCGAN: A regional fusion factor-based union gradient and contrast generative adversarial network for infrared and visible image fusion. *Journal of Modern Optics*, 70(1): 52-68. <https://doi.org/10.1080/09500340.2023.2174358>
- [9] Fu, C., Wu, X., Hu, Y., Huang, H., He, R. (2021). DVG-face: Dual variational generation for heterogeneous face recognition. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 44(6): 2938-2952. <https://doi.org/10.1109/TPAMI.2021.3052549>
- [10] Xiao, Z., Gao, X., Fu, C., Dong, Y., et al. (2021). Improving transferability of adversarial patches on face recognition with generative models. In 2021 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR), Nashville, TN, USA, pp. 11840-11849. <https://doi.org/10.1109/CVPR46437.2021.01167>
- [11] Shang, L., Huang, M., Shi, W., Liu, Y., et al. (2023). Improving training and inference of face recognition models via random temperature scaling. *Proceedings of the AAAI Conference on Artificial Intelligence*, 37(12): 15082-15090. <https://doi.org/10.1609/aaai.v37i12.26760>
- [12] Hariri, W. (2022). Efficient masked face recognition method during the Covid-19 pandemic. *Signal, Image and Video Processing*, 16(3): 605-612. <https://doi.org/10.1007/s11760-021-02050-w>
- [13] Qiu, H., Gong, D., Li, Z., Liu, W., Tao, D. (2021). End2end occluded face recognition by masking corrupted features. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 44(10): 6939-6952. <https://doi.org/10.1109/TPAMI.2021.3098962>
- [14] Zhang, L., Sun, L., Yu, L., Dong, X., et al. (2021). ARFace: Attention-aware and regularization for face recognition with reinforcement learning. *IEEE Transactions on Biometrics, Behavior, and Identity Science*, 4(1): 30-42. <https://doi.org/10.1109/TBIOM.2021.3104014>
- [15] Terhörst, P., Kolf, J.N., Huber, M., Kirchbuchner, F., et al. (2021). A comprehensive study on face recognition biases beyond demographics. *IEEE Transactions on Technology and Society*, 3(1): 16-30. <https://doi.org/10.1109/TTS.2021.3111823>
- [16] Wang, X. (2023). A fuzzy neural network-based automatic fault diagnosis method for permanent magnet synchronous generators. *Mathematical Biosciences and Engineering*, 20(5): 8933-8953. <https://doi.org/10.3934/mbe.2023392>
- [17] Luo, N., Yu, H., You, Z., Li, Y., et al. (2023). Fuzzy logic and neural network-based risk assessment model for import and export enterprises: A review. *Journal of Data Science and Intelligent Systems*, 1(1): 2-11. <https://doi.org/10.47852/bonviewJDSIS32021078>
- [18] Kure, H.I., Islam, S., Ghazanfar, M., Raza, A., Pasha, M. (2022). Asset criticality and risk prediction for an effective cybersecurity risk management of cyber-physical system. *Neural Computing and Applications*, 34(1): 493-514. <https://doi.org/10.1007/s00521-021-06400-0>
- [19] Beke, A., Kumbasar, T. (2022). More than accuracy: A composite learning framework for interval type-2 fuzzy logic systems. *IEEE Transactions on Fuzzy Systems*,

- 31(3): 734-744.
<https://doi.org/10.1109/TFUZZ.2022.3188920>
- [20] Tseng, M.L., Jeng, S.Y., Lin, C.W., Lim, M.K. (2021). Recycled construction and demolition waste material: A cost–benefit analysis under uncertainty. *Management of Environmental Quality: An International Journal*, 32(3): 665-680. <https://doi.org/10.1108/MEQ-08-2020-0175>
- [21] Yu, Q., Song, J.Y., Yu, X.H., Cheng, K., Chen, G. (2022). To solve the problems of combat mission predictions based on multi-instance genetic fuzzy systems. *The Journal of Supercomputing*, 78(12): 14626-14647. <https://doi.org/10.1007/s11227-022-04388-5>
- [22] Zhang, T., Zhan, J., Shi, J., Xin, J., Zheng, N. (2023). Human-like decision-making of autonomous vehicles in dynamic traffic scenarios. *IEEE/CAA Journal of Automatica Sinica*, 10(10): 1905-1917. <https://doi.org/10.1109/JAS.2023.123696>