



iJRASET

International Journal For Research in
Applied Science and Engineering Technology



INTERNATIONAL JOURNAL FOR RESEARCH

IN APPLIED SCIENCE & ENGINEERING TECHNOLOGY

Volume: 14 Issue: VII Month of publication: July 2026

DOI: <https://doi.org/10.22214/ijraset.2026.84467>

www.ijraset.com

Call:  08813907089

E-mail ID: ijraset@gmail.com

Enhanced Deep Convolutional Generative Adversarial Network for High-Fidelity Human Face Synthesis

K. Narasimha Reddy¹, A. Renuka²

¹M-Tech Student, Dept. of CSE, Universal College of Engineering and Technology (Autonomous), Guntur, Andhra Pradesh, India

²Assistant Professor, Dept. of CSE, Universal College of Engineering and Technology (Autonomous), Guntur, Andhra Pradesh, India

Abstract: *Generative Adversarial Networks (GANs) are a potent tool for generating photorealistic synthetic images. In this study, we aim to build and test a Deep Convolutional Generative Adversarial Network (DCGAN) that can produce realistic portraits of people's faces. One network learns to transform random noise vectors into realistic face photographs, while the other learns to differentiate between real and fraudulent photos; both networks are trained on the CelebA dataset. The proposed design uses convolutional layers in the Discriminator and transposed convolutional (deconvolutional) layers in the Generator to improve image quality and training stability. Binary cross-entropy loss, the Adam optimizer, and suitable data normalization are some of the strategies used to ensure effective learning and better convergence. The results prove that DCGANs are capable of creating realistic facial representations. To evaluate the model's efficacy, we use both qualitative and quantitative measures, such as Frechet Inception Distance (FID). FID measures how similar the distributions of the real and generated pictures are to one another. The results of this study demonstrate that DCGANs are very promising for applications in the entertainment, virtual reality, data augmentation, and AI-powered creative content production industries.*

Keywords: *Generative, Adversarial, Networks, Deep Convolutional Neural Network, Visual Inspection.*

I. INTRODUCTION

Generative Adversarial Networks (GANs) have revolutionized picture synthesis in recent years, opening up a world of possibilities for the production of varied and realistic visual material in a variety of disciplines. Among the many GAN designs, Deep Convolutional GANs (DCGANs) stand out for the way they learn complicated spatial patterns and produce stable, high-quality images by integrating convolutional neural networks in the discriminator and generator. Digital content production, VR, gaming, and AI-assisted design systems are just a few of the many areas that might benefit from the realistic human face pictures generated by DCGANs, which is the primary goal of this research. The model is taught to recognise faces using the extensive CelebA dataset. As part of its training process, the generator network learns to create convincing fake faces out of random noise vectors, while the discriminator network enhances the quality of the output by differentiating between the two. For more lifelike picture creation, DCGAN design makes use of convolutional layers, which improve its capacity to grasp spatial hierarchies and fine-grained face characteristics. Although DCGAN models have many benefits, they may also experience problems including unstable training behaviour and mode collapse. These issues can impact the quality and variety of the output. In order to circumvent these restrictions, the suggested system makes use of efficient preprocessing methods, meticulously chosen loss functions, and enhanced training approaches. In sum, the results show that DCGANs can make realistic human faces, which is a big deal in the creative industries and for cutting-edge scientific research.

A. Problem Statement

In this project, we will use a Deep Convolutional Generative Adversarial Network (DCGAN) architecture to create realistic and high-quality photographs of human faces. Using the training dataset's spatial information and fine-grained face details, the model improves the picture synthesis process with the help of convolutional layers. The system is also built to deal with typical problems that come up during GAN training, such mode collapse and instability during convergence. The suggested method seeks to enhance the produced face pictures' quality, variety, and realism by tackling these issues.

B. Challenges in DCGAN

- 1) **Training Instability:** The adversarial nature of the DCGAN's generator and discriminator makes training it an unstable process. Training might become imbalanced as both networks strive to surpass one other. The generator has difficulty learning significant patterns when the discriminator is too strong. On the other hand, when the generator takes center stage, the discriminator stops becoming better. Oscillations or improper convergence during training might result from this mismatch.
- 2) **Difficulty in Convergence:** Doughnut generative adversarial networks (DCGANs) deviate from the conventional wisdom on how to minimise loss. It is challenging to ascertain whether the model has undergone sufficient training due to the lack of a guaranteed global optimal point. Because of this, keeping track of training progress and knowing when to quit becomes more difficult.
- 3) **Limited Control Over Output:** The capacity to change the kind of outputs provided by standard DCGAN models is restricted. Generating class-based pictures or particular face features is challenging due to their absence of conditional procedures. Because of this, precise regulation of the synthesis process is hindered.
- 4) **Computational Cost:** When training, DCGANs need a lot of processing power. They may be demanding on resources because to their reliance on powerful GPUs, significant memory consumption, and lengthy training durations.

C. Proposed Work Aim

One of the main components of the suggested setup is a Deep Convolutional Generative Adversarial Network (DCGAN), which improves upon the performance of regular GANs by adding transposed and convolutional layers. The model's capacity to learn spatial hierarchies is enhanced by this design, leading to the generation of more visually cohesive and realistic representations of human faces.

D. Objectives

Main Purpose:

- Create a deep convolutional neural network (DCGAN) model that can take in random noise and produce human face pictures that seem lifelike.

Additional Goals:

- To guarantee stable and efficient training, it is necessary to preprocess the dataset adequately.
- To get the most out of training and reduce the amount of visible artefacts in the final output.
- Using performance indicators like the Frechet Inception Distance (FID) to assess the quality of produced pictures.

II. LITERATURE SURVEY

- 1) Goodfellow et al. – Generative Adversarial Networks (GANs) (2014) Ian Goodfellow et al. introduced the Generative Adversarial Network (GAN), a revolutionary deep learning framework that fundamentally changed the field of image synthesis. The proposed architecture consists of two neural networks: a Generator, which creates synthetic images from random noise, and a Discriminator, which distinguishes between real and generated images. These two networks compete in a minimax game, allowing the generator to progressively learn the underlying data distribution and produce increasingly realistic images. Unlike traditional generative models, GANs do not require explicit probability density estimation, making them highly effective for generating high-quality synthetic data. This work laid the theoretical foundation for numerous GAN variants and applications in computer vision, image synthesis, medical imaging, and multimedia. However, the original GAN suffered from issues such as unstable training and mode collapse, motivating extensive research into improved GAN architectures and optimization strategies.
- 2) Radford et al. – Unsupervised Representation Learning with Deep Convolutional GANs (DCGAN) (2016)- Radford et al. significantly improved the original GAN by introducing the Deep Convolutional Generative Adversarial Network (DCGAN). Instead of fully connected layers, DCGAN employs convolutional and transposed convolutional layers, enabling the network to learn hierarchical image features automatically. Batch normalization and ReLU activation functions were incorporated to stabilize training and improve convergence. The discriminator uses Leaky ReLU activation to enhance feature extraction. DCGAN demonstrated remarkable improvements in generating realistic human faces, bedrooms, and natural scenes. Moreover, the learned feature representations proved useful for unsupervised image classification and feature learning. This architecture became the standard baseline for most face generation applications due to its simplicity, stability, and ability to produce visually appealing images with relatively small datasets.

- 3) Karras et al. – Progressive Growing of GANs (ProGAN) (2018) To address the challenge of generating high-resolution images, Karras et al. proposed Progressive Growing of GANs (ProGAN). Instead of training the complete high-resolution network from the beginning, the generator and discriminator start with low-resolution images (typically 4×4 pixels). New convolutional layers are gradually added throughout training, allowing the network to learn coarse features before focusing on finer image details. This progressive training strategy significantly improves convergence, reduces instability, and enables the generation of highly realistic images with resolutions up to 1024×1024 pixels. ProGAN achieved outstanding performance on celebrity face datasets and became one of the first GAN architectures capable of generating photo-realistic facial images suitable for practical applications.
- 4) Karras et al. – StyleGAN: A Style-Based Generator Architecture (2019)-Building upon ProGAN, Karras et al. introduced StyleGAN, one of the most influential GAN architectures for realistic face generation. Rather than directly feeding the latent vector into the generator, StyleGAN introduces a mapping network that transforms the latent vector into an intermediate latent space. Adaptive Instance Normalization (AdaIN) is then used to control image styles at different layers of the generator. This enables independent manipulation of facial attributes such as hairstyle, facial expression, age, lighting, pose, and identity while preserving image realism. StyleGAN also introduced stochastic noise inputs that generate natural variations such as skin texture and hair details. The model produces highly photorealistic faces and offers exceptional control over image synthesis, making it widely adopted in image editing, digital art, virtual humans, and face generation research.
- 5) Brock et al. – BigGAN: Large Scale GAN Training for High Fidelity Natural Image Synthesis (2019) Brock et al. proposed BigGAN, which focuses on improving image quality through large-scale training. The authors demonstrated that increasing model capacity, dataset size, and batch size substantially enhances image fidelity. BigGAN incorporates spectral normalization, orthogonal regularization, and class-conditional batch normalization to stabilize training while preserving diversity in generated images. Trained on the ImageNet dataset, BigGAN achieved state-of-the-art performance in terms of both visual quality and quantitative evaluation metrics. Although originally designed for natural image synthesis, its training methodology inspired subsequent high-resolution face generation models by demonstrating the importance of large datasets and scalable architectures.
- 6) Wang et al. – GAN-Based Human Face Image Synthesis (2024) Wang et al. presented a comprehensive study on GAN-based human face image synthesis, reviewing several modern GAN architectures developed specifically for realistic facial image generation. The study compares classical GAN, DCGAN, Conditional GAN, StyleGAN, Progressive GAN, and attention-based GAN models with respect to image quality, diversity, training stability, computational complexity, and controllability. The authors highlight that advanced GAN models effectively capture facial attributes including age, gender, facial expression, hairstyle, and pose while preserving identity consistency. The review also discusses major challenges such as mode collapse, limited training datasets, computational cost, and ethical concerns associated with synthetic face generation. The study concludes that combining attention mechanisms, conditional learning, and high-resolution architectures provides the best performance for realistic face synthesis.
- 7) Zhang et al. – Self-Attention Generative Adversarial Networks (SAGAN) (2019)-Zhang et al. proposed the Self-Attention Generative Adversarial Network (SAGAN) to overcome the limited receptive field of convolutional layers. Traditional convolutional GANs primarily capture local image features, making it difficult to maintain global consistency. SAGAN introduces self-attention modules that allow distant image regions to interact during feature learning, enabling the model to capture long-range spatial dependencies. Spectral normalization is applied to both the generator and discriminator, resulting in more stable adversarial training. Experimental results demonstrate significant improvements in image realism, object consistency, and fine-detail preservation. For facial image generation, SAGAN effectively models global facial structures, including eye alignment, facial symmetry, and texture consistency, leading to higher-quality synthesized images than conventional convolutional GANs.

The evolution of GAN-based image synthesis demonstrates a clear progression in addressing the limitations of earlier models. The original GAN introduced adversarial learning for realistic image generation. DCGAN enhanced image quality using convolutional architectures and stable training techniques. Progressive GAN enabled high-resolution image generation through gradual network expansion, while StyleGAN introduced style-based latent representations for controllable face synthesis. These developments collectively establish Deep Convolutional Generative Adversarial Networks (DCGANs) as a strong foundation for advanced face synthesis and provide valuable directions for future research in high-resolution, identity-preserving, and controllable image generation

III. PROPOSED METHODOLOGY

An improved version of the classic GAN design that incorporates transposed convolutional layers is the foundation of the suggested system, the Deep Convolutional Generative Adversarial Network (DCGAN). With this enhancement, the model can now learn spatial hierarchies more efficiently, leading to higher-quality, more realistic representations of human faces.

The System's Proposed Features

- 1) Generator: A generator network can take random noise vectors and turn them into pictures of faces that seem quite human. The input noise is gradually upsampled into a structured visual representation using transposed convolutional layers.
- 2) Stereotype: Differentiating between produced and genuine pictures is the job of the discriminator. To determine whether a picture is genuine or not, it uses convolutional layers to get detailed spatial information.
- 3) The Procedure for Training: During training, the model is subjected to alternating updates from both networks, a technique known as adversarial learning. Binary Cross-Entropy (BCE) loss is used as a learning aid by the system. During training, the Adam optimiser is used with the parameters $\alpha = 0.0002$ and $\beta_1 = 0.5$ to guarantee efficient and steady convergence.

A. Algorithm

Input:

- Training dataset of real face images (X)
- Noise distribution (Z) (e.g., Gaussian or Uniform)
- Number of training epochs (E)
- Batch size (B)
- Learning rate (α)

Output:

- Trained Generator (G) capable of generating realistic face images.

B. Algorithm Steps

- 1) Initialize the Generator (G) network with random weights.
- 2) Initialize the Discriminator (D) network with random weights.
- 3) Define the loss function.
- 4) For each epoch $e = 1$ to E :

For each batch of real images from dataset X:

a) Train the Discriminator

1. Sample a random noise vector z .
2. Generate fake images:

$$x_{fake} = G(z)$$

3. Compute discriminator outputs:

$$D_{real} = D(x)$$

$$D_{fake} = D(x_{fake})$$

4. Compute discriminator loss:

$$Loss_D = -[\log(D_{real}) + \log(1 - D_{fake})]$$

5. Update discriminator parameters using gradient descent:

$$\theta_D = \theta_D - \alpha \nabla (Loss_D)$$

b) Train the Generator

1. Sample another random noise vector z .
2. Generate fake images:

$$x_{fake} = G(z)$$

3. Compute generator loss:

$$Loss_G = -\log(D(G(z)))$$

4. Update generator parameters using gradient descent:

$$\theta_G = \theta_G - \alpha \nabla (Loss_G)$$

5. Monitor training progress:

- Generate sample images using $G(z)$.
- Evaluate model performance (optional).

6. Repeat until all epochs are completed.

Return:

Trained Generator model G .

DCGAN Optimization Objective:

The DCGAN optimizes the following minimax objective:

$$G_{\min} D_{\max} V(D, G) = E_{x \sim p_{data}(x)} [\log D(x)] + E_{z \sim p_z(z)} [\log (1 - D(G(z)))]$$

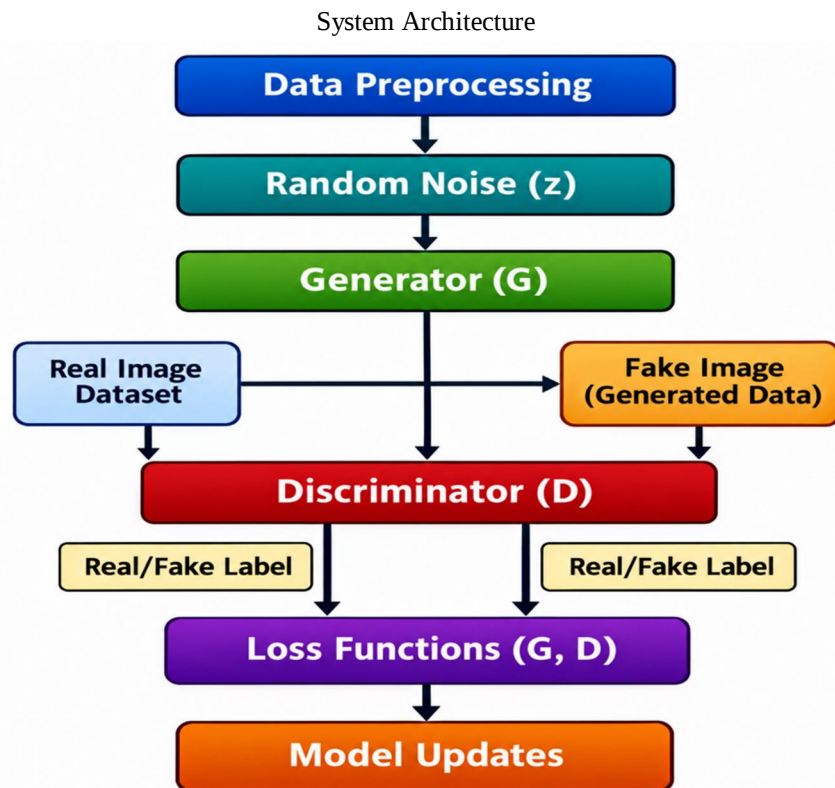


Fig 1: System Architecture

IV. RESULTS AND DISCUSSIONS

A. Data Collection and Preprocessing

1) Data Collection

A Deep Convolutional Generative Adversarial Network (DCGAN) cannot be built without first collecting data. The quantity, variety, and quality of the training dataset significantly impact the model's efficacy. More realistic and high-quality images are generated when the model learns rich feature representations from a dataset that has been carefully selected.

Data Origins:

Accessible databases: Popular benchmark datasets like CIFAR-10, MNIST, and CelebA (CelebFaces Attributes Dataset) are usually a part of the dataset that trains a Deep Convolutional Generative Adversarial Network (DCGAN).

The CIFAR-10 dataset is ideal for testing general picture generating ability as it contains varied natural images from several item categories.

The first testing and validation of generative models often use MNIST, which is composed of handwritten digit pictures, because of its simplicity and structured nature.

For training models aimed at realistic human face synthesis, CelebA (CelebFaces Attributes Dataset) is an excellent choice due to its extensive collection of annotated photos of human faces.

Training and evaluating DCGAN models is made possible by these datasets, which provide a good mix of simplicity, variety, and complexity.

Private data sets:

Depending on the field of use, datasets may be compiled from a variety of sources. To guarantee accuracy and reliability, data is often acquired via controlled acquisition systems, digital cameras, or sensors. The collection of massive picture datasets from internet archives and open-source platforms is another common usage of web scraping methods.

Model performance in certain domains is enhanced using domain-specific datasets for specialised applications. Some examples of such datasets include those used in medicine (including X-ray and MRI images), agriculture (including pictures of plant leaves for disease diagnosis), and remote sensing and geographic analysis (using satellite photos). The generalisability and resilience of deep learning models, such as DCGANs, are enhanced by using such varied data sources.

Needed Information:

To train deep learning models like DCGAN effectively, a dataset needs a big enough picture size to learn complicated patterns. More data means less chance of overfitting and greater generalisation by the model.

Another important thing is that the dataset should be quite diverse. This means that there should be a lot of different forms, textures, lighting, positions, and backdrops in the dataset. The model is able to learn more accurate and generalised data representations because to this variety.

All photos must adhere to the same format, which includes consistent preparation and uniform resolution. To make training go smoothly and to boost the model's stability and efficiency, make sure the picture size and structure are consistent.

2) *Cleaning Up Your Data*

Preprocessing the dataset is necessary for consistent training and enhanced model performance before it is fed into a DCGAN. To ensure the model learns well and produces high-quality synthetic pictures, proper preprocessing is essential.

Procedures for Preprocessing

Step One: Image Resizing

A set resolution, such 64×64 or 128×128 pixels, is applied to all scaled photos. For consistent processing, convolutional neural networks need input dimensions that are uniform, hence this stage is essential.

B. *Standardisation*

Normalisation methods are used to scale pixel values from the range of 0 to 255 to the range of -1 to 1. The numerical stability is improved and the model may converge quicker during training with this change.

C. *Cleaning the Data*

We exclude low-quality, irrelevant, or corrupted photos from the collection. This guarantees that training is based on just relevant and valuable data, which improves learning results.

D. *Supplementing Data (Optional)*

Augmentation methods including rotating, flipping, cropping, and zooming are used to boost dataset variety. The generator's capacity for generalisation is enhanced by these modifications.

E. Making Batches

Datasets are often broken into smaller batches, with each batch containing 32 or 64 samples. This enables for more efficient calculation and more consistent gradient updates during training.

F. Merging Data Sets

Before training, the sequence of photos is randomly chosen. This improves training stability generally and prevents the model from picking up on any random sequence patterns.

Thirdly, Preprocessing Output Here is the dataset after preprocessing:

To prepare the dataset for training the DCGAN model, preparation involves cleaning, normalising, and resizing the dataset to a standard format.

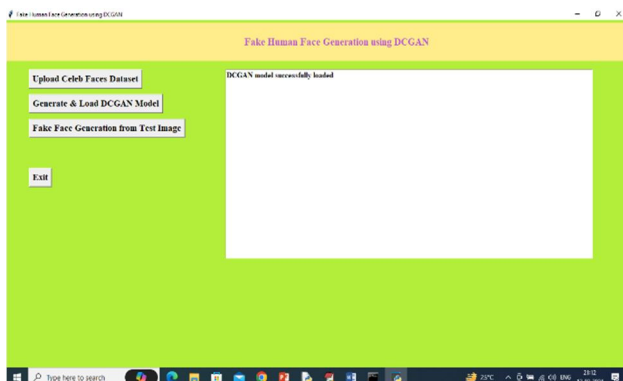


Fig 2: DCGAN Model Successfully Loaded

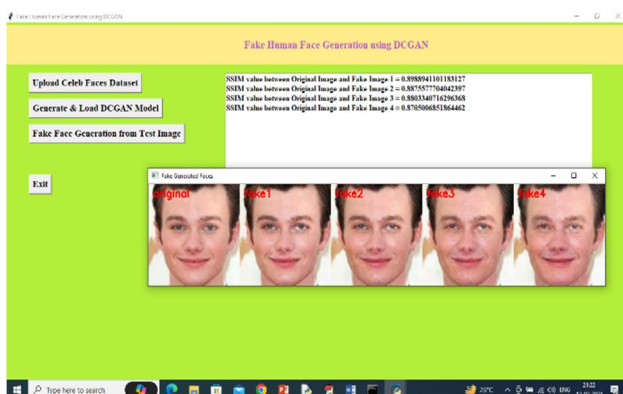
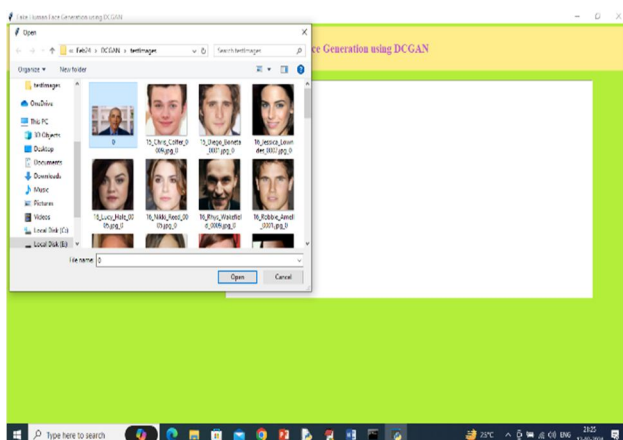


Fig.3 Generated Fake Face Images with SSIM Evaluation



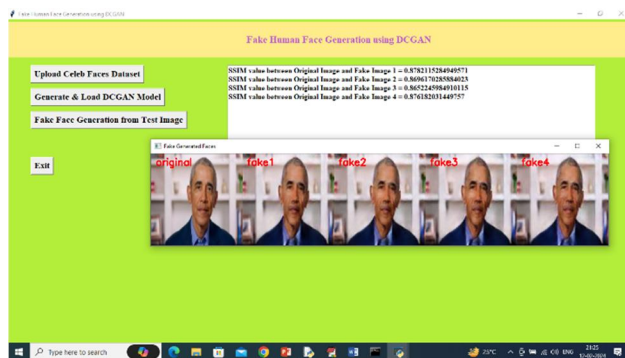


Fig 4: Obama Image Generation Results

V. CONCLUSION

The study demonstrates the efficacy of deep learning approaches in creating lifelike artificial visuals by creating synthetic human face pictures using Deep Convolutional Generative Adversarial Networks (DCGAN). The system can generate high-quality face pictures from noise vectors that have been randomly sampled by using a DCGAN architecture that consists of a Discriminator and a Generator. The core learning mechanism of the model is the adversarial training process, in which the Generator tries to trick the Discriminator while the latter learns to differentiate between actual and false pictures. This study concludes that DCGANs can produce realistic pictures of human faces and sheds light on the difficulties of adversarial training. Virtual reality, data augmentation, digital content production, and simulation-based settings are just a few of the promising uses for the created system. For even more realistic and detailed image creation in the future, it may be necessary to investigate more complex GAN structures, boost training stability, and improve picture resolution.

A. Scope for Future Work

While the present DCGAN-based system does a good job at creating synthetic human face pictures, there are a number of ways it might be improved in the future to make it even better. The goal of these improvements is to make the model stronger and more efficient by addressing current issues with picture quality, scalability, training stability, and assessment accuracy. Improving the produced pictures' resolution may be a future goal for developers looking to make their products seem more lifelike. To further increase control over produced characteristics and boost visual quality, more complex GAN designs like StyleGAN or versions of conditional GAN may be used. Better optimisation tactics, stabilised training methods, and enhanced loss functions are some of the solutions that may be used to fix problems like mode collapse and training instability. The system's capabilities may be enhanced with the addition of conditional generation and multi-modal inputs, enabling more controlled and diversified picture creation. The model would be more realistic for implementation in real-world applications if it included real-time processing and optimised computing efficiency. Think about the moral implications, such how to utilise synthetic faces responsibly and how to stop people from abusing the material that is created. In the long run, these upgrades have the potential to greatly expand the system's capabilities, making it more applicable in a variety of fields including entertainment, virtual reality, and AI-driven content production, as well as boosting the quality of the created photos.

REFERENCES

- [1] I. J. Goodfellow, J. Pouget-Abadie, M. Mirza, B. Xu, D. Warde-Farley, S. Ozair, A. Courville, and Y. Bengio, "Generative Adversarial Nets," in Proceedings of the 28th Annual Conference on Neural Information Processing Systems (NeurIPS), Montreal, QC, Canada, 2014, pp. 2672–2680.
- [2] A. Radford, L. Metz, and S. Chintala, "Unsupervised Representation Learning with Deep Convolutional Generative Adversarial Networks," in Proceedings of the International Conference on Learning Representations (ICLR), San Juan, Puerto Rico, 2016.
- [3] T. Salimans, I. Goodfellow, W. Zaremba, V. Cheung, A. Radford, and X. Chen, "Improved Techniques for Training GANs," in Proceedings of the 30th Annual Conference on Neural Information Processing Systems (NeurIPS), Barcelona, Spain, 2016, pp. 2234–2242.
- [4] M. Arjovsky, S. Chintala, and L. Bottou, "Wasserstein GAN," in Proceedings of the 34th International Conference on Machine Learning (ICML), Sydney, Australia, 2017, pp. 214–223.
- [5] T. Karras, T. Aila, S. Laine, and J. Lehtinen, "Progressive Growing of GANs for Improved Quality, Stability, and Variation," in Proceedings of the International Conference on Learning Representations (ICLR), Vancouver, BC, Canada, 2018.
- [6] M. Mirza and S. Osindero, "Conditional Generative Adversarial Nets," arXiv preprint arXiv:1411.1784, 2014.
- [7] X. Dong and Y. Zhang, "Face Image Synthesis with Deep Convolutional Generative Adversarial Networks," in Proceedings of the International Conference on Artificial Intelligence and Computer Engineering (ICAICE), Beijing, China, 2020, pp. 297–301.

- [8] K. He, X. Zhang, S. Ren, and J. Sun, "Deep Residual Learning for Image Recognition," in Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR), Las Vegas, NV, USA, 2016, pp. 770–778.
- [9] A. Brock, J. Donahue, and K. Simonyan, "Large Scale GAN Training for High Fidelity Natural Image Synthesis," in Proceedings of the International Conference on Learning Representations (ICLR), New Orleans, LA, USA, 2019.
- [10] J.-Y. Zhu, T. Park, P. Isola, and A. A. Efros, "Unpaired Image-to-Image Translation Using Cycle-Consistent Adversarial Networks," in Proceedings of the IEEE International Conference on Computer Vision (ICCV), Venice, Italy, 2017, pp. 2223–2232.
- [11] M.-Y. Liu and O. Tuzel, "Coupled Generative Adversarial Networks," in Proceedings of the 30th Annual Conference on Neural Information Processing Systems (NeurIPS), Barcelona, Spain, 2016, pp. 469–477.
- [12] J. Johnson, A. Alahi, and L. Fei-Fei, "Perceptual Losses for Real-Time Style Transfer and Super-Resolution," in Proceedings of the 14th European Conference on Computer Vision (ECCV), Amsterdam, The Netherlands, 2016, pp. 694–711.
- [13] I. Goodfellow, Y. Bengio, and A. Courville, Deep Learning. Cambridge, MA, USA: MIT Press, 2016.
- [14] R. Szeliski, Computer Vision: Algorithms and Applications, 2nd ed. Cham, Switzerland: Springer, 2022.
- [15] D. P. Kingma and M. Welling, "Auto-Encoding Variational Bayes," in Proceedings of the International Conference on Learning Representations (ICLR), Banff, AB, Canada, 2014.
- [16] A. Krizhevsky, I. Sutskever, and G. E. Hinton, "ImageNet Classification with Deep Convolutional Neural Networks," in Proceedings of the 25th Annual Conference on Neural Information Processing Systems (NeurIPS), Lake Tahoe, NV, USA, 2012, pp. 1097–1105.
- [17] O. Ronneberger, P. Fischer, and T. Brox, "U-Net: Convolutional Networks for Biomedical Image Segmentation," in Proceedings of the International Conference on Medical Image Computing and Computer-Assisted Intervention (MICCAI), Munich, Germany, 2015, pp. 234–241.
- [18] A. Dosovitskiy et al., "An Image Is Worth 16×16 Words: Transformers for Image Recognition at Scale," in Proceedings of the International Conference on Learning Representations (ICLR), Vienna, Austria, 2021.
- [19] P. Isola, J.-Y. Zhu, T. Zhou, and A. A. Efros, "Image-to-Image Translation with Conditional Adversarial Networks," in Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR), Honolulu, HI, USA, 2017, pp. 1125–1134.
- [20] T. Karras, S. Laine, and T. Aila, "A Style-Based Generator Architecture for Generative Adversarial Networks," in Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR), Long Beach, CA, USA, 2019, pp. 4401–4410.

FIRST AUTHORS:

K NARASIMHA REDDY pursuing his M. Tech in Computer Science and Engineering in Universal College of Engineering and Technology.

SECOND AUTHOR:

Mrs. A. RENUKA, M. Tech, she is currently working as Assistant Professor of CSE in Universal College of Engineering and Technology



10.22214/IJRASET



45.98



IMPACT FACTOR:
7.129



IMPACT FACTOR:
7.429



INTERNATIONAL JOURNAL FOR RESEARCH

IN APPLIED SCIENCE & ENGINEERING TECHNOLOGY

Call : 08813907089  (24*7 Support on Whatsapp)