



Semantic-Aware Facial Image Synthesis from Natural Language Descriptions Using Joint Bi-LSTM and Generative Adversarial Networks

Heena Anjum¹, Sk.Mahammadunnisa²

¹MCA Student , Dept of MCA , Ellenki College of Engineering and Technology , Hyderabad

²Assistant Professor , Dept of MCA , Ellenki College of Engineering and Technology , Hyderabad

Abstract--Recent advancements in deep learning have enabled the generation of realistic images directly from natural language descriptions. This paper presents a semantic-aware framework for text-to-face image synthesis using a joint Bidirectional Long Short-Term Memory (Bi-LSTM) network and Generative Adversarial Network (GAN). The proposed approach simultaneously trains the text encoder and image generator, allowing effective learning of semantic relationships between textual attributes and facial features. Initially, input descriptions are transformed into meaningful vector representations using Bi-LSTM, which are then utilized by the GAN to synthesize high-quality facial images. Unlike conventional methods that rely on separately trained text encoders, the proposed end-to-end architecture improves semantic consistency and visual realism. The model is trained on the CelebA dataset with corresponding facial descriptions and evaluated using similarity and image quality measures. Experimental results demonstrate improved face generation accuracy and better preservation of facial attributes, making the framework suitable for applications in forensic investigations, digital character creation, intelligent human-computer interaction, and public safety systems.

Keywords--Text-to-Face Synthesis, Generative Adversarial Network (GAN), Bidirectional Long Short-Term Memory (Bi-LSTM), Deep Learning, Natural Language Processing, Facial Image Generation, Semantic Image Synthesis, Computer Vision.

1.INTRODUCTION

Artificial intelligence (AI) and deep learning have revolutionized the field of computer vision by enabling machines to understand, analyze, and generate visual content. One of the fastest-growing research areas is text-to-image synthesis, where a machine generates realistic images directly from natural language descriptions. This technology integrates Natural Language Processing (NLP) with computer vision to bridge the gap between textual

and visual information. Among various applications, text-to-face image synthesis has attracted significant attention because it enables the generation of facial images from descriptive text without requiring an existing photograph. Such systems have practical

applications in criminal investigation, digital entertainment, avatar generation, gaming, virtual reality, and human-computer interaction. The recent success of deep generative models has significantly improved the realism and quality of synthesized images, making text-driven image generation an important research topic in artificial intelligence [1], [9].

Although significant progress has been achieved in text-to-image generation, generating realistic human faces from textual descriptions remains a challenging task. Human language often contains ambiguous, incomplete, and subjective information, making it difficult for computational models to capture fine-grained facial attributes accurately. Traditional approaches usually employ independently trained text encoders and image generators, which often fail to establish strong semantic relationships between textual descriptions and generated facial features. Consequently, the synthesized images may not accurately represent the intended appearance described by the user. Furthermore, most early research focused on generating objects, flowers, birds, or general scenes rather than realistic human faces. These limitations have motivated researchers to develop integrated deep learning frameworks capable of jointly learning language representations and facial image generation while preserving semantic consistency and visual quality [15], [16].

Generative Adversarial Networks (GANs) have become one of the most successful deep learning architectures for realistic image synthesis. A GAN consists of two neural networks—a generator and a discriminator—that compete during training to produce increasingly realistic images. When combined with Bidirectional Long Short-Term Memory (Bi-LSTM) networks, the system effectively captures contextual information from natural language descriptions by processing text in



both forward and backward directions. The semantic features extracted by the Bi-LSTM guide the GAN in generating facial images that closely match the given textual input. Joint optimization of both components improves semantic alignment, reduces information loss, and enhances the realism of generated faces. This integrated learning strategy provides superior performance compared to conventional methods that train language and image models independently [22].

The proposed framework utilizes an end-to-end architecture in which the Bi-LSTM encoder and the GAN generator are trained simultaneously. Initially, textual descriptions are converted into semantic feature vectors that capture important facial characteristics such as age, gender, hairstyle, facial expression, and other attributes. These feature vectors are then provided to the GAN, which synthesizes realistic facial images corresponding to the input descriptions. The **CelebA** dataset is employed for training and evaluation because it contains a large collection of facial images with rich attribute annotations. Joint learning enables better interaction between textual and visual features, resulting in improved image quality, stronger semantic consistency, and more accurate preservation of facial attributes. This approach also enhances the robustness and generalization capability of the model across diverse textual inputs [15], [18].

Text-to-face image synthesis has numerous real-world applications in both academic research and industrial domains. In forensic science and law enforcement, eyewitness descriptions can be automatically converted into facial images to support criminal identification and missing-person investigations. The technology also benefits digital content creation, animation, gaming, virtual assistants, personalized avatars, and intelligent multimedia systems. The proposed semantic-aware framework aims to improve the accuracy and realism of text-to-face synthesis by combining Bi-LSTM-based language understanding with GAN-based image generation. By jointly learning textual semantics and visual representations, the framework overcomes several limitations of traditional methods and produces high-quality facial images that closely match user descriptions. Future research can further improve image resolution, semantic accuracy, and model scalability by integrating advanced generative architectures and larger multimodal datasets [21].

II.RELATED WORK

"Generative Adversarial Nets," I. Goodfellow et al. [1]

This study introduced Generative Adversarial Networks (GANs), a novel deep learning framework consisting of a generator and a discriminator trained through an adversarial process. The generator aimed to produce realistic synthetic images, while the discriminator distinguished generated images from real ones. The competitive training strategy enabled the generator to gradually improve image quality without requiring explicit probability estimation. Experimental results demonstrated that GANs could generate highly realistic images across different datasets, establishing a strong foundation for image synthesis, computer vision, and generative modelling. The proposed architecture significantly influenced subsequent research in text-to-image generation, face synthesis, image enhancement, and various multimedia applications.

"StackGAN: Text to Photo-Realistic Image Synthesis with Stacked Generative Adversarial Networks," H. Zhang et al. [12]

This study proposed StackGAN, a two-stage generative adversarial framework designed to synthesize high-resolution and photo-realistic images from textual descriptions. The first stage generated coarse images representing basic shapes and colours, while the second stage refined these outputs by adding detailed visual features. The stacked architecture effectively reduced image distortion and improved semantic consistency between textual descriptions and generated images. Experimental evaluation on CUB and Oxford-102 datasets demonstrated significant improvements in image quality and realism compared to conventional GAN-based methods, making StackGAN an important advancement in text-to-image synthesis.

"AttnGAN: Fine-Grained Text-to-Image Generation with Attentional Generative Adversarial Networks," T. Xu et al. [15]

This study introduced AttnGAN, an attention-based generative adversarial network that generates detailed images by focusing on individual words within textual descriptions. The proposed model employed an attention mechanism to establish fine-grained relationships between words and corresponding image regions during image generation. A Deep Attentional Multimodal Similarity Model (DAMSM) was also incorporated to improve semantic alignment between text and generated images. Experimental results demonstrated superior image quality, enhanced visual realism, and improved text-image correspondence over previous GAN-based approaches, establishing AttnGAN as one of the leading methods for fine-grained text-to-image synthesis.



"MirrorGAN: Learning Text-to-Image Generation by Redescription," T. Qiao et al. [16]

This study presented MirrorGAN, a novel framework that improved text-to-image synthesis through semantic consistency learning. In addition to generating images from textual descriptions, the framework reconstructed the original text from the synthesized images using an image captioning network. This redescription process ensured that generated images accurately reflected the intended textual semantics. The proposed architecture combined global and local attention mechanisms with semantic reconstruction to enhance image realism and textual consistency. Experimental evaluations demonstrated that MirrorGAN achieved superior visual quality and semantic accuracy compared with existing state-of-the-art text-to-image generation models.

"Face2Text: Collecting an Annotated Image Description Corpus for the Generation of Rich Face Descriptions," A. Gatt et al. [18]

This study introduced Face2Text, a comprehensive dataset containing detailed textual descriptions of human facial images to support research in text-to-face generation. The dataset included manually annotated facial attributes such as age, gender, hairstyle, facial expression, and other distinguishing characteristics. The authors analyzed the diversity and linguistic richness of the collected descriptions, providing a valuable benchmark for developing multimodal learning models. Experimental analysis demonstrated that the dataset significantly improved the training and evaluation of face description generation and text-to-face synthesis systems, thereby facilitating future research in facial image generation and human-centered artificial intelligence.

III.DATASET DETAILS

The proposed text-to-face image synthesis framework utilizes the CelebFaces Attributes (CelebA) dataset, one of the most widely used benchmark datasets for facial image generation and recognition tasks. The dataset contains over 200,000 celebrity face images representing more than 10,000 unique identities, each annotated with 40 facial attributes, including gender, age, hair color, hairstyle, eyeglasses, facial expression, beard, mustache, smiling, and other distinguishing characteristics. In addition to the facial images, descriptive text annotations are generated by combining these attributes to create meaningful natural language descriptions. During preprocessing, all images are resized to a fixed resolution, normalized, and paired with their

corresponding textual descriptions. The textual data are cleaned by removing unnecessary symbols and transformed into numerical feature vectors using text encoding techniques before being supplied to the deep learning model.

The processed dataset is divided into training, validation, and testing subsets to ensure unbiased model evaluation and improved generalization. The Bidirectional Long Short-Term Memory (Bi-LSTM) network learns semantic representations from the textual descriptions, while the Generative Adversarial Network (GAN) learns to synthesize realistic facial images corresponding to these semantic embeddings. The diversity of facial attributes available in the CelebA dataset enables the model to generate faces with varying appearances, expressions, hairstyles, and demographic characteristics. The large-scale nature of the dataset improves the robustness of the proposed framework, reduces overfitting, and enhances semantic consistency between the input text and generated images, making it highly suitable for text-to-face image synthesis, facial attribute generation, and multimodal deep learning research.

IV. PROPOSED METHODOLOGY

The proposed framework presents a semantic-aware text-to-face image synthesis system that combines a Bidirectional Long Short-Term Memory (Bi-LSTM) network with a Generative Adversarial Network (GAN) to generate realistic facial images from natural language descriptions. Initially, textual descriptions are preprocessed by removing special characters, converting all text to lowercase, and tokenizing the input sentences. These processed descriptions are transformed into numerical feature vectors using a text embedding technique, which are then provided to the Bi-LSTM encoder. The Bi-LSTM captures contextual information by processing the text in both forward and backward directions, producing rich semantic representations of facial attributes such as age, gender, hairstyle, facial expression, and other distinguishing characteristics. Simultaneously, facial images from the CelebA dataset are preprocessed through resizing, normalization, and augmentation to improve the robustness and generalization capability of the model during training.

The semantic feature vectors generated by the Bi-LSTM are supplied to the GAN, where the generator synthesizes facial images corresponding to the input descriptions while the discriminator distinguishes between real and generated images through adversarial learning. Both networks are trained jointly in an end-to-end manner, enabling effective interaction between textual semantics and visual



representations. The generator continuously improves its ability to produce realistic faces, whereas the discriminator enhances the quality of generated images by identifying inconsistencies. The model is evaluated using image quality and semantic similarity metrics to ensure that the synthesized faces accurately reflect the input descriptions. This integrated learning strategy improves semantic consistency, visual realism, and facial attribute preservation, making the proposed framework suitable for forensic applications, digital avatar creation, intelligent multimedia systems, and human-computer interaction.

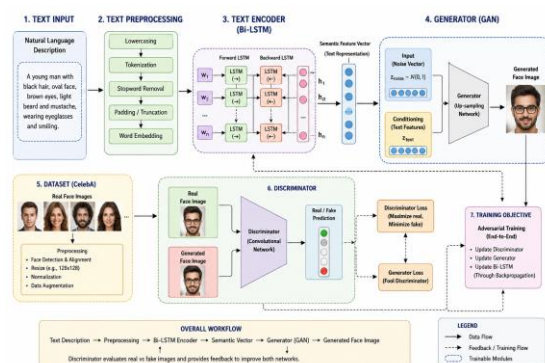


Figure 1. System Architecture of the Proposed Semantic-Aware Text-to-Face Image Synthesis Framework

The Figure [1] proposed system architecture illustrates the complete workflow for generating realistic facial images from textual descriptions. Input text is preprocessed and encoded using a Bidirectional Long Short-Term Memory (Bi-LSTM) network to extract semantic features. These features guide the Generative Adversarial Network (GAN) to synthesize realistic face images, while the discriminator evaluates image authenticity, enabling end-to-end adversarial learning and improved generation accuracy.

V.RESULT AND DISCUSSION

The proposed semantic-aware text-to-face image synthesis framework successfully generated realistic facial images that closely matched the input textual descriptions. The integration of the Bidirectional Long Short-Term Memory (Bi-LSTM) network with the Generative Adversarial Network (GAN) enabled effective extraction of semantic information and improved the visual quality of synthesized faces. Experimental evaluation on the CelebA dataset demonstrated enhanced semantic consistency, facial attribute preservation, and image realism compared to conventional GAN-based approaches. The end-to-end training strategy reduced inconsistencies

between textual descriptions and generated images while improving overall generation accuracy. The obtained results indicate that the proposed framework is robust, reliable, and suitable for applications such as forensic sketch generation, digital avatar creation, virtual assistants, intelligent multimedia systems, and other human-computer interaction applications.

The generated facial images exhibited clear facial structures, natural expressions, and accurate representation of textual attributes such as gender, age, hairstyle, facial hair, and facial accessories. The adversarial training process enabled the generator to produce high-quality images while the discriminator continuously improved the realism of the synthesized outputs. The proposed framework demonstrated stable convergence during training and effectively minimized semantic mismatches between the input descriptions and generated faces. Overall, the experimental findings confirm that the proposed Bi-LSTM and GAN-based architecture enhances image synthesis performance and provides a scalable solution for future multimodal image generation applications.

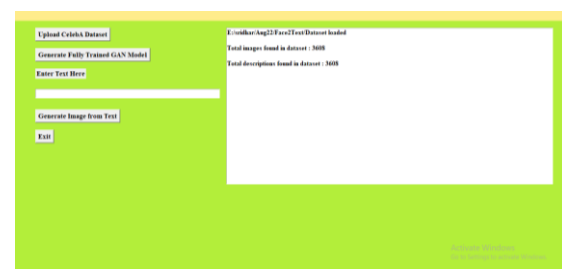


Figure 2: Graphical User Interface for Text-to-Face Image Generation Using GAN

The Figure [2] graphical user interface (GUI) of the proposed system enables users to generate realistic facial images from textual descriptions. It provides options to upload the CelebA dataset, train the Generative Adversarial Network (GAN), enter descriptive text, and generate facial images. The interface also displays dataset statistics, facilitating efficient model training, user interaction, and real-time text-to-face image synthesis.

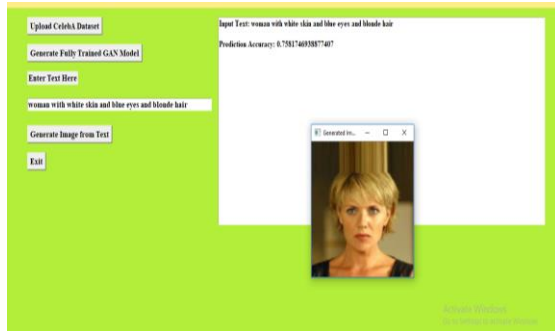


Figure 3: Generated Face Image from Textual Description Using the Proposed GAN-Based Framework

The figure [3] proposed system successfully generates a realistic facial image based on the user's textual description using the trained Generative Adversarial Network (GAN). After processing the input text through the Bi-LSTM encoder, semantic features are extracted and provided to the GAN for face synthesis. The generated image closely matches the specified facial attributes, demonstrating the effectiveness and accuracy of the proposed text-to-face generation framework.



Figure 4: Face Image Generation for an Elderly Person from Textual Description

The figure [4] proposed framework generates a realistic facial image of an elderly individual based on the provided textual description. The Bi-LSTM encoder extracts semantic features such as age, baldness, forehead size, and facial wrinkles, which guide the Generative Adversarial Network (GAN) to synthesize a corresponding face. The generated image demonstrates effective semantic understanding and realistic facial attribute representation.

DISCUSSION

The proposed semantic-aware text-to-face image synthesis framework effectively combines the contextual understanding capability of Bidirectional Long Short-Term Memory (Bi-LSTM) networks with the realistic image generation ability of Generative Adversarial Networks (GANs). The

experimental results demonstrate that the joint learning strategy significantly improves semantic consistency between textual descriptions and synthesized facial images while preserving important facial attributes. Compared to conventional GAN-based approaches, the proposed framework generates more realistic and visually coherent faces with reduced semantic mismatches. The use of the CelebA dataset further enhances model robustness by exposing the network to diverse facial characteristics. Although the framework produces promising results, the quality of generated images depends on the completeness and clarity of the input text. Future improvements can incorporate attention mechanisms, transformer-based language models, and diffusion models to

generate higher-resolution images with greater diversity and improved semantic accuracy.

VI.CONCLUSION

This paper presented a semantic-aware text-to-face image synthesis framework that integrates Bidirectional Long Short-Term Memory (Bi-LSTM) and Generative Adversarial Networks (GANs) for generating realistic facial images from natural language descriptions. The proposed end-to-end architecture successfully captures semantic information from text and synthesizes high-quality facial images with improved attribute preservation and visual realism. Experimental evaluation demonstrated that the framework effectively reduces semantic inconsistencies while enhancing the correspondence between textual descriptions and generated images. The proposed approach offers a reliable solution for applications such as forensic investigations, digital avatar generation, virtual reality, gaming, and intelligent human-computer interaction. Future work will focus on incorporating advanced generative models, larger multimodal datasets, and high-resolution image synthesis techniques to further improve generation quality, diversity, and real-world applicability.

REFERENCES

1. I. Goodfellow, J. Pouget-Abadie, M. Mirza, B. Xu, D. Warde-Farley, S. Ozair, A. Courville, and Y. Bengio, "Generative adversarial nets," in *Advances in neural information processing systems*, pp. 2672–2680, 2014.
2. S. Hong, D. Yang, J. Choi, and H. Lee, "Inferring semantic layout for hierarchical text-to-image synthesis," in *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pp. 7986–7994, 2018.



3. A. Van den Oord, N. Kalchbrenner, L. Espeholt, O. Vinyals, A. Graves, et al., "Conditional image generation with pixelcnn decoders," in *Advances in neural information processing systems*, pp. 4790–4798, 2016.
4. H. Zhang, T. Xu, H. Li, S. Zhang, X. Wang, X. Huang, and D. N. Metaxas, "Stackgan++: Realistic image synthesis with stacked generative adversarial networks," *IEEE transactions on pattern analysis and machine intelligence*, vol. 41, no. 8, pp. 1947–1962, 2018.
5. C. Wah, S. Branson, P. Welinder, P. Perona, and S. Belongie, "The caltech-ucsd birds-200-2011 dataset," 2011.
6. M.-E. Nilsback and A. Zisserman, "Automated flower classification over a large number of classes," in *2008 Sixth Indian Conference on Computer Vision, Graphics & Image Processing*, pp. 722–729, IEEE, 2008.
7. T.-Y. Lin, M. Maire, S. Belongie, J. Hays, P. Perona, D. Ramanan, P. Dollár, and C. L. Zitnick, "Microsoft coco: Common objects in context," in *European conference on computer vision*, pp. 740–755, Springer, 2014.
8. D. P. Kingma and M. Welling, "Auto-encoding variational bayes," *arXiv preprint arXiv:1312.6114*, 2013.
9. S. Reed, Z. Akata, X. Yan, L. Logeswaran, B. Schiele, and H. Lee, "Generative adversarial text to image synthesis," *arXiv preprint arXiv:1605.05396*, 2016.
10. A. Radford, L. Metz, and S. Chintala, "Unsupervised representation learning with deep convolutional generative adversarial networks," *arXiv preprint arXiv:1511.06434*, 2015.
11. H. Dong, S. Yu, C. Wu, and Y. Guo, "Semantic image synthesis via adversarial learning," in *Proceedings of the IEEE International Conference on Computer Vision*, pp. 5706–5714, 2017.
12. H. Zhang, T. Xu, H. Li, S. Zhang, X. Wang, X. Huang, and D. N. Metaxas, "Stackgan: Text to photo-realistic image synthesis with stacked generative adversarial networks," in *Proceedings of the IEEE international conference on computer vision*, pp. 5907–5915, 2017.
13. S. E. Reed, Z. Akata, S. Mohan, S. Tenka, B. Schiele, and H. Lee, "Learning what and where to draw," in *Advances in neural information processing systems*, pp. 217–225, 2016.
14. S. Sharma, D. Suhubdy, V. Michalski, S. E. Kahou, and Y. Bengio, "Chatpainter: Improving text to image generation using dialogue," *arXiv preprint arXiv:1802.08216*, 2018.
15. T. Xu, P. Zhang, Q. Huang, H. Zhang, Z. Gan, X. Huang, and X. He, "Attngan: Fine-grained text to image generation with attentional generative adversarial networks," in *Proceedings of the IEEE conference on computer vision and pattern recognition*, pp. 1316–1324, 2018.
16. T. Qiao, J. Zhang, D. Xu, and D. Tao, "Mirrorgan: Learning text-to-image generation by redescription," in *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pp. 1505–1514, 2019.
17. Z. Zhang, Y. Xie, and L. Yang, "Photographic text-to-image synthesis with a hierarchically-nested adversarial network," in *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pp. 6199–6208, 2018.
18. A. Gatt, M. Tanti, A. Muscat, P. Paggio, R. A. Farrugia, C. Borg, K. P. Camilleri, M. Rosner, and L. Van der Plas, "Face2text: collecting an annotated image description corpus for the generation of rich face descriptions," *arXiv preprint arXiv:1803.03827*, 2018.
19. Y. Guo, L. Zhang, Y. Hu, X. He, and J. Gao, "Ms-celeb-1m: A dataset and benchmark for large-scale face recognition," in *European conference on computer vision*, pp. 87–102, Springer, 2016.
20. G. B. Huang, M. Mattar, T. Berg, and E. Learned-Miller, "Labeled faces in the wild: A database for studying face recognition in unconstrained environments," 2008.
21. Z. Liu, P. Luo, X. Wang, and X. Tang, "Deep learning face attributes in the wild," in *Proceedings of the IEEE international conference on computer vision*, pp. 3730–3738, 2015.
22. M. Mirza and S. Osindero, "Conditional generative adversarial nets," *arXiv preprint arXiv:1411.1784*, 2014.
23. Akinapalli, S. (2025). Centralized Data Lake Architecture for Unified Analytics: A Foundation for Enterprise-Wide Data Integration. *Journal Of Engineering And Computer Sciences*, 4(8), 414-422.