

Face Aging Using Generative Adversarial Network

P.V.S. Lakshmi Jagadamba^{1,*}, Amrutha Gudimetla², Srija Katru²,
Pujitha Asapu², Harshitha Kandregula²

Abstract

This project addresses the challenge of predicting how a person may look in the future or how they appeared in the past using a single photograph. While existing methods mainly focus on altering texture, they often neglect changes in head shape that naturally occur during the aging process, limiting their effectiveness, especially when applied to images of children. To tackle this issue, a novel approach is introduced that employs a multi-domain image-to-image generative adversarial network architecture. This innovative framework captures a continuous bi-directional adversarial aging process in its learned latent space. By training our network on the FFHQ dataset, meticulously annotated for age, gender, and semantic segmentation, establishing fixed age classes as reference points to approximate seamless age transformation, our model is capable of generating full head portraits spanning ages 0 to 70 years from a single input photograph, encompassing both texture and head shape modifications. Through extensive experimentation across diverse datasets, it showcases substantial enhancements over existing techniques, underscoring the efficacy and versatility of our proposed methodology.

Keywords: Head shape, single input photograph, bi-directional, semantic segmentation, efficacy

INTRODUCTION

Transforming a person's appearance to simulate aging while maintaining their identity poses a significant challenge, especially when dealing with substantial age gaps. Existing methods often struggle with simultaneously altering both head shape and facial texture, particularly when attempting to create a continuous spectrum of age transformations spanning from infancy to older adulthood, rather than just binary transitions.

While some approaches focus on minor age differences or primarily on adult-to-elderly progression, they typically prioritize texture modifications over shape alterations [1–6]. Techniques like those

*Author for Correspondence

P.V.S. Lakshmi Jagadamba
E-mail: drpvsljagadamba@gvpcew.ac.in

¹Professor and Head, Department of Computer Science and Engineering, Gayatri Vidya Parishad College of Engineering (Autonomous) (GVPCE), Visakhapatnam, Andhra Pradesh, India

²Student, Department of Computer Science and Engineering, Gayatri Vidya Parishad College of Engineering (Autonomous) (GVPCE), Visakhapatnam, Andhra Pradesh, India

Received Date: October 25, 2024
Accepted Date: November 07, 2024
Published Date: January 08, 2025

Citation: P.V.S. Lakshmi Jagadamba, Amrutha Gudimetla, Srija Katru, Pujitha Asapu, Harshitha Kandregula. Face Aging Using Generative Adversarial Network. Journal of Image Processing & Pattern Recognition Progress. 2025; 12(1): 41–52p.

employed by Kemelmacher-Shlizerman *et al.*, offer substantial age transformations but are limited to cropped face areas and lack the capability for backward age prediction [7]. The concept of Face Aging, encompassing the synthesis of ages from 0 to 70 years, presents a theoretical challenge due to the continuous nature of time. To address this, we propose a representing age with a fixed number of anchor classes in a multi-domain transfer framework, facilitating the learning of geometric transformations across different age groups while covering the entire age spectrum. Our approach involves a novel multi-domain image-to-image conditional GAN architecture, where an identity encoder extracts features from input images, age domains are represented by unique distributions, and a mapping network facilitates continuous age

transformations in a unified latent space. While disentanglement-based domain transfer methods such as MUNIT and FUNIT excel at learning shape and texture deformation for distinct categories like cats and dogs, they face limitations when applied to age transformation due to scalability issues and potential attribute interference. Similarly, multi-domain transfer algorithms like STARGAN and STGAN struggle with age transformation tasks due to the high correlation between age domains. Our proposed algorithm overcomes these challenges by effectively transforming both shape and texture across a wide age range while preserving the individual's identity and characteristics.

Dataset limitations, particularly the scarcity of baby and children's images in existing face-aging datasets, prompted us to label the FFHQ dataset for gender and age attributes, along with extracting additional facial features such as semantic maps and head pose angles. Through qualitative and quantitative evaluations, we demonstrate the superiority of our method over existing aging algorithms, multi-domain transfer approaches, and latent space traversal methods. Our contributions include enabling comprehensive shape and texture transformations for aging synthesis, introducing a novel multi-domain image-to-image translations GAN architecture, and providing the labelled FFHQ dataset to the research community. Ethical considerations and potential biases inherent in such methods are acknowledged and addressed in our supplementary Ethics and Bias statement.

LITERATURE SURVEY

Antipov *et al.* propose "Face Ageing with Conditional Generative Adversarial Networks," which addresses the difficulty of face ageing using conditional generative adversarial networks (GANs) [1]. Their solution makes use of the IMDB-Wiki cleansed dataset and presents a new methodology for optimising GAN latent vectors while maintaining the original person's identity. Despite their best efforts, significant negatives include the possibility of identity loss and the inherent blurriness of generated images, demanding more refinement for more accurate and identity-preserving results.

He *et al.* investigate the distribution of ageing determinants and ageing trends across ages and people [2]. Their research, presented at the 2019 IEEE International Conference on Computer Vision (ICCV), looks into ways to generalise ageing effects and patterns in facial recognition tasks. Their effort aims to increase the resilience and efficiency of age-related facial analysis systems.

Kemelmacher-Shlizerman *et al.* presented "Age Progression Cognizant of Illumination" at the IEEE Conference on Pattern Recognition and Computer Vision (CVPR) in 2014 [7]. Their research focuses on age progression techniques that account for changes in illumination in order to improve the realism and accuracy of generated aged faces. Their technique, which takes into account illumination effects in the ageing process, adds to more reliable age progression models for use in computer vision and facial analysis.

Duong *et al.* present their "Temporal Non-Volume Preserving Method to Facial Age-Progression and Age-Invariant Face Recognition" at the 2017 IEEE International Conference on Computer Vision (ICCV) [3]. They use a temporal non-volume preserving strategy to address facial age progression and age-invariant face recognition. They hope to improve the durability and accuracy of facial recognition systems across different ages, adding to advances in biometric security and identity verification technology.

Wang *et al.* [5] explored recurrent face ageing and reported their findings at the IEEE Conference on Computer Vision and Pattern Recognition Proceedings in 2016. Cui *et al.* investigated recurrent models for face ageing analysis [4]. They hope to increase our understanding and modelling of temporal ageing trends in facial images by investigating recurrent architectures, which will contribute to advances in facial analysis and biometric applications.

Tang *et al.* presented "Personalized age progression with bi-level aging dictionary learning" at the IEEE Conference on Pattern Recognition and Computer Vision (CVPR) in 2017 [6]. Their research

focuses on using conditional adversarial autoencoders for age progression and regression applications. Their goal is to improve the accuracy and realism of age progression and regression models by combining autoencoder architecture with conditional adversarial training. Their approach helps to develop facial age analysis and synthesis techniques, allowing for further applications in age-related research and biometrics.

PROPOSED SYSTEM

The suggested AI-powered solution makes use of the StyleGAN deep learning architecture to produce short video clips that depict an individual's ageing process [8]. The method begins with a high-resolution frontal image and preprocesses it by resizing, cropping, and normalisation before putting it into the StyleGAN model. StyleGAN, the architecture's basic component that includes a conditional generator and discriminator, has been pre-trained on a large dataset of faces with age labels. The conditional generator, which is made up of an identity encoder, a mapping network, and a decoder, coordinates transitions across age groups. By recording age and identity separately, the approach assures that a person's appearance varies with age, but their identity does not [9].

The suggested method represents each age group with separate predefined distributions, allowing for the sampling of target age codes [10]. The mapping network transforms these codes into a unified age latent space, allowing for ongoing age alterations. At the same time, the identity encoder extracts feature from the input image. The decoder generates realistic face images by combining the mapping network output and the target age latent vector via modulated convolutions inspired by StyleGAN2. During training, an age encoder matches real and produced images to the distributions of their corresponding age classes [11]. When ages are not represented, the algorithm estimates age latent codes for neighbouring classes and applies linear interpolation. This system demonstrates AI's capacity to realistically replicate human ageing, with intriguing applications in entertainment and digital art and aging research (Figure 1).

METHODOLOGY

The StyleGAN face ageing project employs a methodical strategy to provide people with an engaging age transformation experience. To preserve privacy and data security standards, it starts with a secure user authentication procedure. After completing the authentication process, customers upload a single photo of the person they want to age simulate, which starts the system's processing pipeline [12].

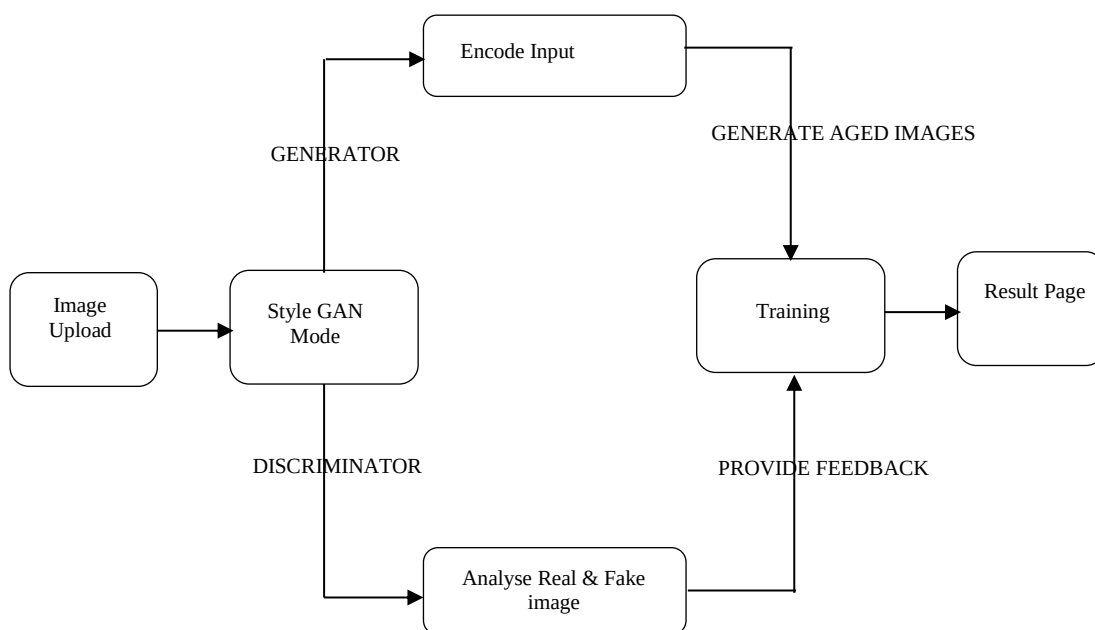


Figure 1. Workflow of the proposed system.

The input image undergoes preprocessing procedures by the system once it has been uploaded. These processes, which include resizing, cropping, and normalisation, maximise compliance with the StyleGAN model, the system's core, that has already been trained. This model is skilled at producing ageing effects that are convincing on faces.

The input image and the intended target age are fed into the StyleGAN model after post-processing. By utilising its potential, the model creates a series of face pictures that represent the person at different points in their life, accurately representing ageing [13]. Users can see how the submitted image changes visually at various age milestones thanks to this process.

After a brief video clip is created, the produced facial images are subjected to interpolation algorithms to create a smooth transition between age phases. Users may see the person's ageing process in this video clip, which shows them from their young age to an older age in a visually appealing manner [14]. Users can view and download the generated video of the age transition through the system's user-friendly interface (Figure 2).

FACE AGE GENERATION

The main goal is to create an algorithm that can learn the entire range of head shape deformation and appearance changes across a wide age range [15]. The standard method to deal with this problem would be supervised learning, ideally. Nevertheless, as ageing is a continual process, it requires a significant number of aligned image pairs representing the same person at different ages to account for all possible conversions. Sadly, there are no large-scale datasets that currently exist that capture ageing changes over prolonged periods of time, let alone a lifetime. Furthermore, supervised training becomes extremely challenging because available small-scale datasets such as FGNET present participants in a variety of positions, surroundings, and lighting situations. Therefore, we leverage recent advances in unpaired image-to-image translation GAN frameworks and resort to adversarial learning [16].

Suggest a novel generative adversarial network design that consists of a single conditional generator and discriminator in order to address this issue. The conditional generator, which consists of three main parts: an identity encoder, a mapping network, and a decoder, takes on the task of handling transitions between age groups.

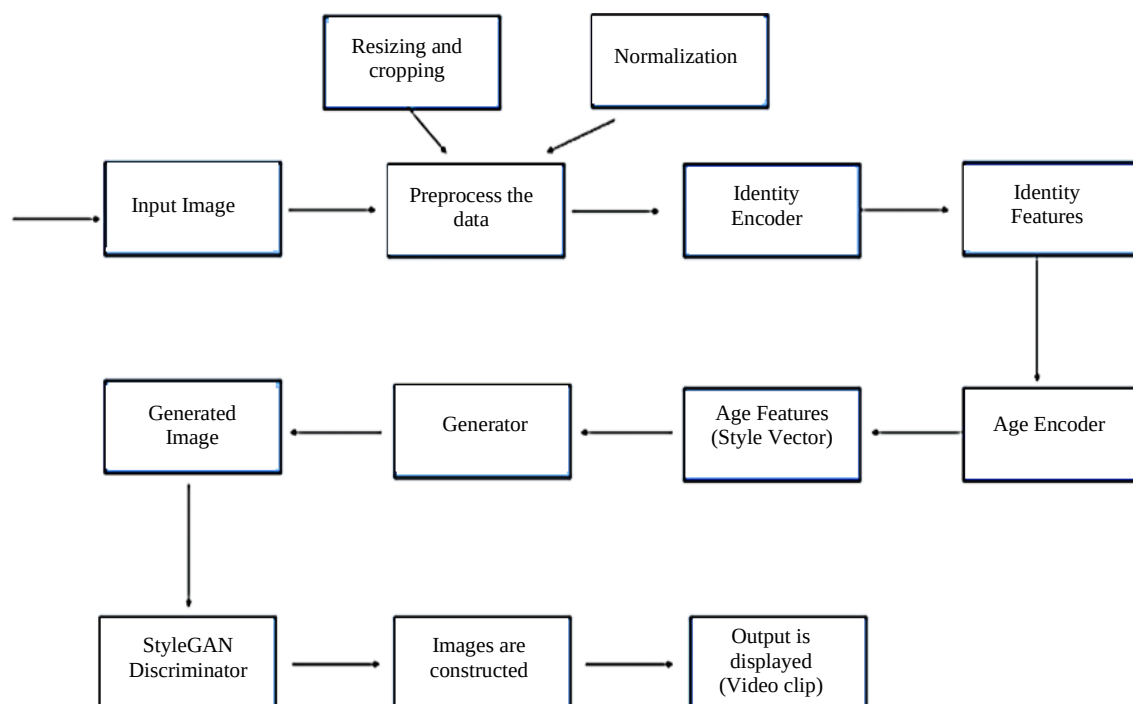


Figure 2. Architecture of proposed model.

We propose that an individual's identity does not change as they age, but their appearance does. As so, we divide identity and age along separate trajectories [17]. There is a specific preset distribution that defines each age group. assign a vector age code to a sample that is taken from the corresponding age group distribution given the target age. Then, this code is put into a mapping network to approximate continuous age transformations by translating it into a learnt unified age latent space.

In order to extract identification features, the identity encoder processes the input image simultaneously. The decoder uses modulated convolutions, as first suggested in StyleGAN2, to mix the output of the mapping network with the target age latent vector and inject it into the identity features. An extra age encoder helps with the training process by enabling the correlation between generated and real images with the established age class distributions [18]. When we have two nearby anchor classes, we compute age latent codes for each one, and then we use linear interpolation to acquire the necessary age code as input for the decoder. This process is used when transformation to an age not represented in our anchor classes is needed (Figure 3).

Framework

Using a facial image, the algorithm creating an output image of the same person at the specified age cluster and a target age cluster as inputs [19]. This generator consists of three separate parts: a decoder, mapping network, and identity encoder. It also includes an age encoder for training. Furthermore, a discriminator is employed to distinguish between authentic and synthesised images.

Pre-processing and Age Input Representation

Using its semantic mask, the input image's background and garment components are eliminated. A vector, $z_i \in Z$, with each member corresponding to an age class, represents the age input space Z .

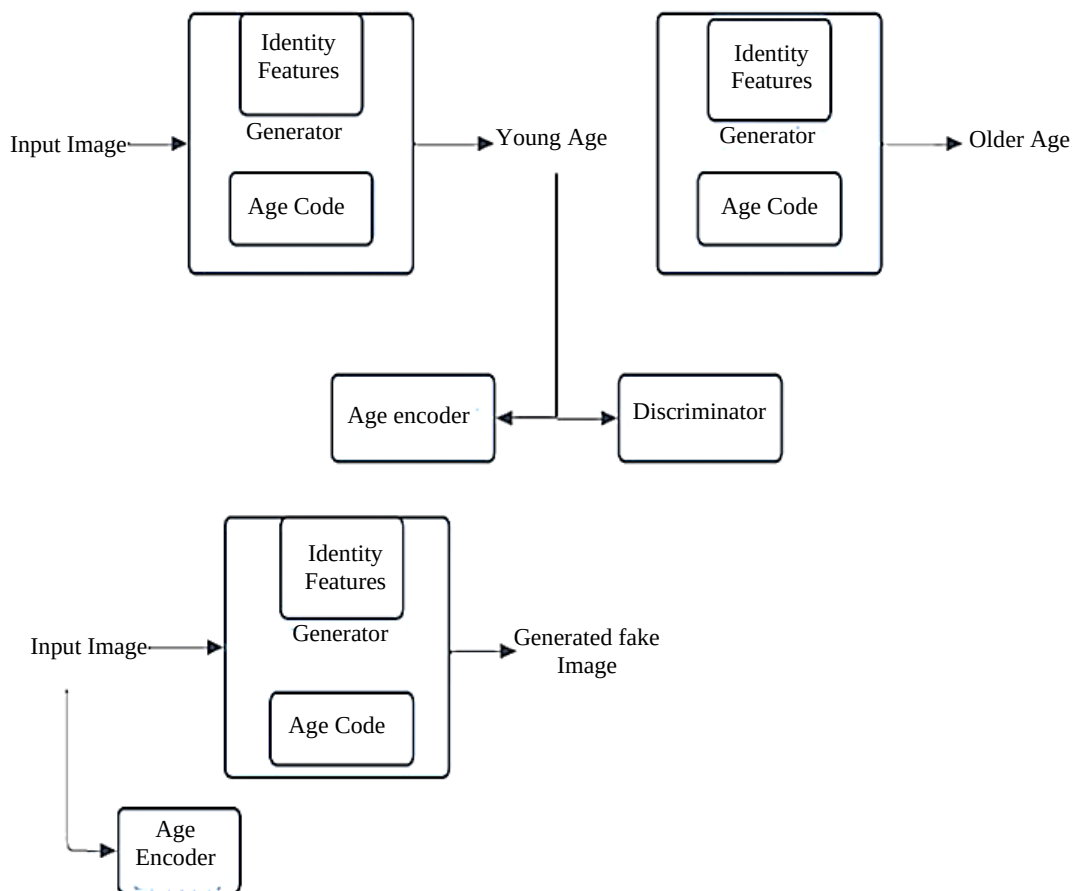


Figure 3. Algorithm Overview.

The target age class is taken into consideration when computing z_i , which is then utilised to generate images at the desired age cluster.

$$Z_i = 1i + v, v \sim N(0, 0.2^2 \cdot I)$$

Identification Encoder

The encoder of identity Eid takes an input picture x and uses $wid = Eid(x)$ to extract the identity features tensor. These features hold details about the general shape of the face and the local structures of the image, both of which are important in creating the same identity. There are two down sampling stages and four residual blocks in the identity encoder.

Mapping Network

An age input vector, z , is embedded into the unified age latent space $Wage$, $wage = M(z)$, via the mapping network $M: Z \rightarrow Wage$, where M is an 8-layer MLP network and $wage$ is a 256-element latent vector. In order to provide seamless transitions and interpolations between age clusters, which are necessary for ongoing age transformations, the mapping network learns an ideal age latent space.

Decoder

The decoder generates an output picture $y = F(wid, wage)$ from an age latent code and identification characteristics. Styled convolution blocks are used to process the Lifespan Age Transformation Synthesis identity features. We swap out the AdaIN normalisation layers with the modulated convolution layers suggested in StyleGAN2 in order to lessen the water droplet artefacts. Additionally, as we have seen, a pixel norm layer comes after each modulated convolution layer, which helps to further reduce these artefacts. In our implementation, we do not include the noise injection. In all, we employ two up sampling styled convolution layers to create an image at its original size and four styled convolution layers to work with the identity code.

From an input image x and an input target age vector z_t to an output image y , the overall generator mapping is as follows:

$$y = G(x, z_t) = F(Eid(x), M(z_t))$$

Age Encoder

In the age vector space Z , the age encoder mandates that the input picture x be mapped to its proper place. It generates an age vector $z_s = Eage(x)$ that matches the image x 's source age cluster, s . To encode the overall appearance, irrespective of identity, the age encoder must obtain additional global data. In order to do this, we use a fully linked layer, global averaging, and four down sampling layers, as well as the architecture of the style encoder from MUNIT to generate an age vector. Keep in mind that inference is not done using the age encoder.

Discriminator

With minibatch standard deviation, we employ the StyleGAN discriminator. In order to differentiate between many classes, we alter the final fully connected layer to include n outputs. We only penalise the i -th output for a real image from class I . Accordingly, given a created image of class j , only the j -th output is penalised.

Training Procedure

For each training iteration, a random selection of a source age cluster (s) and a target age cluster (t) is made in order to ensure diversity in age transitions. To enable variation in training data, a sample of images is subsequently taken from each age group.

Each iteration consists of three forward passes after age cluster selection and picture sampling. In order to complete these passes, an image must be generated at the target age t , reconstructed at the source age s , and a cycle consistency check must be carried out in order to reconstruct the image at the source age from the generated image at t_{gen} .

In order to direct the generator and discriminator's optimisation, the training procedure includes a variety of loss functions. These loss functions include identity feature loss (L_{id}), which preserves an individual's identity as they age, adversarial loss (L_{adv}), self-reconstruction loss (L_{rec}) to enforce identity translation, and age vector loss (L_{age}) to ensure proper embedding of real and generated images into the input age space.

These loss functions are combined to create the overall optimisation function, with each term being weighted by the appropriate coefficients (λ_{rec} , λ_{cyc} , λ_{id} , λ_{age}). In order to improve the generator and discriminator and produce high-quality images with precise age transformations and preserved identities, the training process iteratively minimises this combined loss.

Dataset

FFHQ-Aging, a new facial ageing dataset created from photos in the FFHQ dataset, was created as part of the dataset preparation process [20]. Three judgements were made for each of the 70,000 FFHQ photos, and gender and age cluster annotations were gathered using the Appen crowd-sourcing app. Throughout a person's life, from 0–2 years old to 70+ years old, the age clusters were designed to represent both geometric and appearance variations.

Based on the CelebAMask-HQ dataset training, a DeepLabV3 network was used to extract face semantic maps for every image. The Face++ platform was also used to acquire ocular occlusion ratings, head position angles, and type of glasses. Following the alignment process from a prior study, but using a slightly bigger crop size, 256×256 resolution photographs and semantic maps were produced.

Images were divided into training and testing sets after the dataset was created, and precise pruning criteria were used. We eliminated images with severe head position angles, dark spectacles, high eye occlusion scores, and low gender or age confidence. After additional refinement, the training set was divided into six age groups: 0–2, 3–6, 7–9, 15–19, 30–39, and 50–69 years. After everything was said and done, the training set consisted of 14,232 male and 14,066 female photographs in the end, plus 198 male and 205 female images for testing.

Implementation Details

Two distinct models, one designed for male participants and the other for female subjects, were trained in the implementation details. Four NVIDIA GeForce RTX 2080 Ti GPUs were used for each model's training, with a batch size of 12 samples per iteration. 400 epochs of training were used to optimise the model's parameters and reach convergence. By leveraging the Adam optimizer, we ensure successful gradient descent by setting the parameters $\beta_1=0$, $\beta_1=0$ and $\beta_2=0.999$, $\beta_2=0.999$. Moreover, during training, a constant learning rate of 10^{-3} , 10^{-3} was used.

Every 50 and 100 epochs, the learning rate was systematically decreased by a factor of 0.5 to improve training stability and speed up convergence. By making this dynamic change, the optimisation process was able to converge towards an ideal solution while also adapting to the changing dynamics of the model. For this training regimen, four NVIDIA GeForce RTX 2080 Ti GPUs, which are renowned for their excellent speed and efficiency in deep learning tasks were used as computational resources.

Our training approach optimised the generator's performance by combining the non-saturating adversarial loss with R1 regularisation, in accordance with the concepts of StyleGAN. Additionally, in

order to improve the mapping network, we applied exponential moving averages in conjunction with a 0.01 learning rate reduction to stabilise training dynamics and improve overall model performance. Specific weighting coefficients were carefully assigned to guarantee a balanced influence of different loss components. This allowed each loss term to contribute appropriately to the overall training objective, facilitating both efficient age transformation synthesis and identity preservation.

Hyperparameter Tuning

A key component of optimising model performance in our StyleGAN-based face ageing study is hyperparameter adjustment. The quantity of epochs and the batch size are two important hyperparameters that significantly influence the training dynamics and the calibre of the output.

How many iterations the model goes through when training on the dataset is determined by the number of epochs. An increased epoch count could provide the model with additional chances to identify complex patterns in the data, which could result in better age-transformed face image synthesis. But be careful if sufficient regularisation techniques are not used, overly high epoch counts may increase the likelihood of overfitting.

In the same way, changing the batch size can have a significant effect on the training procedure. A training procedure's stability and convergence are influenced by the batch size, which also influences how many samples are handled in a given training cycle.

Leveraging more data per iteration, a bigger batch size could potentially accelerate convergence and speed up training. On the other hand, in situations where computational resources are few, smaller batch sizes may be advantageous in terms of memory efficiency and convergence stability. To balance training effectiveness and model stability, it is necessary to adjust the batch size, which will ultimately improve the overall performance of the StyleGAN-based face ageing system.

Evaluation

Comparison with Paid Apps

A qualitative comparison is conducted with the FaceApp3 outcomes. FaceApp offers binary face ageing filters to alter a person's age perception. Figure 4 illustrates that even with FaceApp's excellent output image quality, it is primarily restricted to skin texture and is unable to conduct shape alteration.

In order to achieve age transformations, we utilised FaceApp's "old" and "cool old" filters and contrasted the results with our output for individuals aged 50–69 years. Our 15–19 years class is approximately identical to the "young2" filter, which we used for changes to a younger age. In order to illustrate our 2-class outputs, we also present our sites.

Comparison with Age Transformation Methods

We assessed both qualitative and quantitative elements of performance in our comparison analysis using three state-of-the-art age transformation methods: IPCGAN, PyGAN, and S2GAN.

It is trained on FFHQ and tested on CACD, and then we compared PyGAN with S2GAN using the CACD dataset for qualitative evaluation. Our single-generator strategy achieved improved photorealism across many age classes, while PyGAN uses a different generator for each age cluster. When it came to representing wrinkles and facial features as age increased, our algorithm outperformed S2GAN in terms of accuracy, covering a wider range of age changes.

Furthermore, our approach was assessed against IPCGAN models that were trained on the FFHQ-Aging and CACD datasets. We improved upon IPCGAN by retraining it on the FFHQ-Aging dataset (referred to as "IPCGAN-retrained"), which allowed for a more equitable comparison.

In addition, we performed a user survey to measure opinions about PyGAN outcomes in comparison to our method. The study evaluated overall preference, perceived age closeness to target age, and subject identity retention. According to our hypothesis, PyGAN might not be able to achieve large age changes, even though it performs exceptionally well at identity preservation. This work led to additional improvements in age transformation algorithms by offering insightful information about the relative advantages and disadvantages of each approach (Table 1).

RESULT AND ANALYSIS

The Face Ageing Project uses the robust StyleGAN deep learning architecture to produce an engaging visual story of a person's ageing process. The procedure starts with a single input image that has been meticulously pre-processed to provide the best possible quality and lighting. After that, the StyleGAN model, which was trained on a big dataset of faces from a range of age ranges, transforms this image.

In the process, a series of realistic face images depicting the person at various periods of life are produced by the StyleGAN model by manipulation of latent codes. The technology is able to accurately replicate the complex alterations in facial appearance that come with ageing by utilising these latent codes. Each generated image adds to the realism of the scene by capturing minute details like skin texture, wrinkles, and facial shapes.

Table 1. Comparison of Ours vs. PyGAN.

Age range	50–69 PyGAN	50–69 Ours
Same identity	20	13
Age difference	24.5	6.9
Overall better	5	15

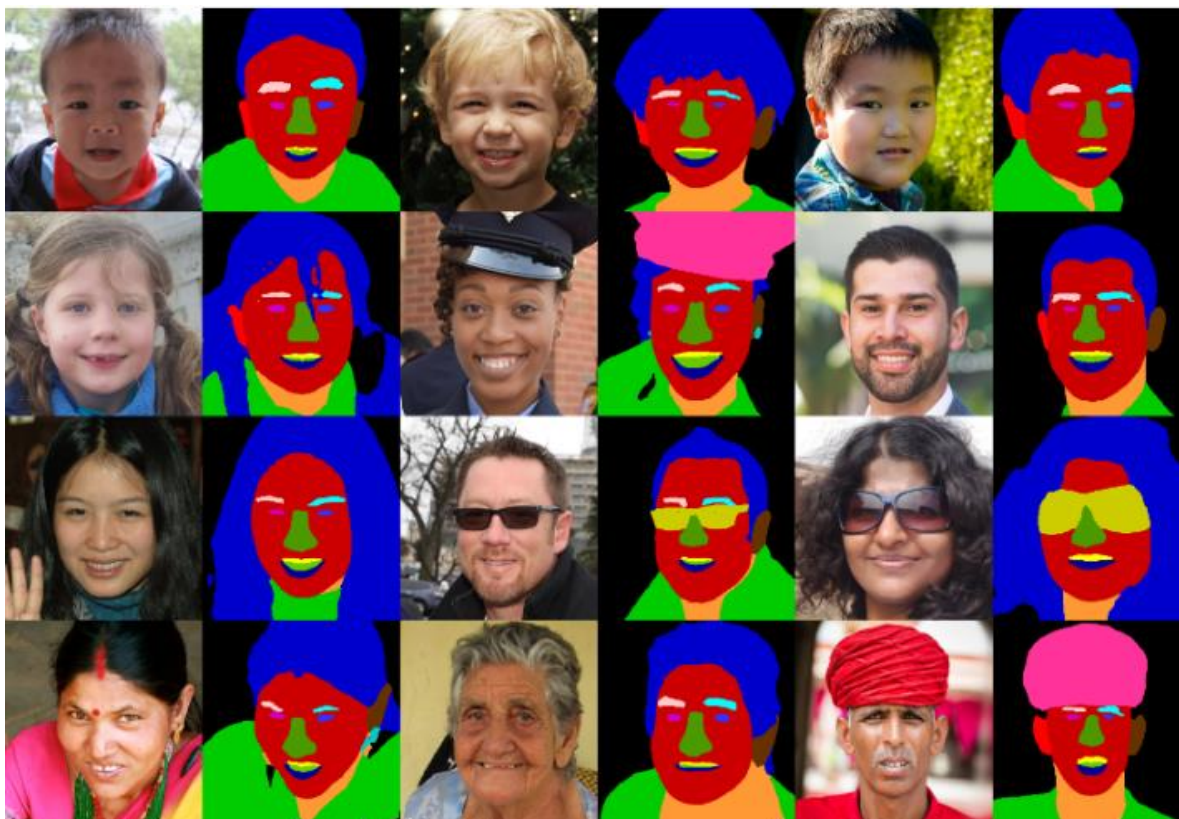


Figure 4. FFHQ-aging dataset sample.

Interpolation techniques are utilised to provide a smooth transition between distinct age phases, guaranteeing a seamless progression from the individual's current appearance to an older age.

A powerful video clip that effectively captures the subject's apparent ageing process is produced as a result of this skilfully combined collection of photos.

The research exhibits the amazing potential of AI to replicate the natural ageing process by synthesising realistic facial alterations and seamlessly integrating age progression.

The final video clip provides an incredibly realistic and detailed depiction of the person's ageing process while also providing insights on the possible effects of ageing on face features (Figures 5–7).

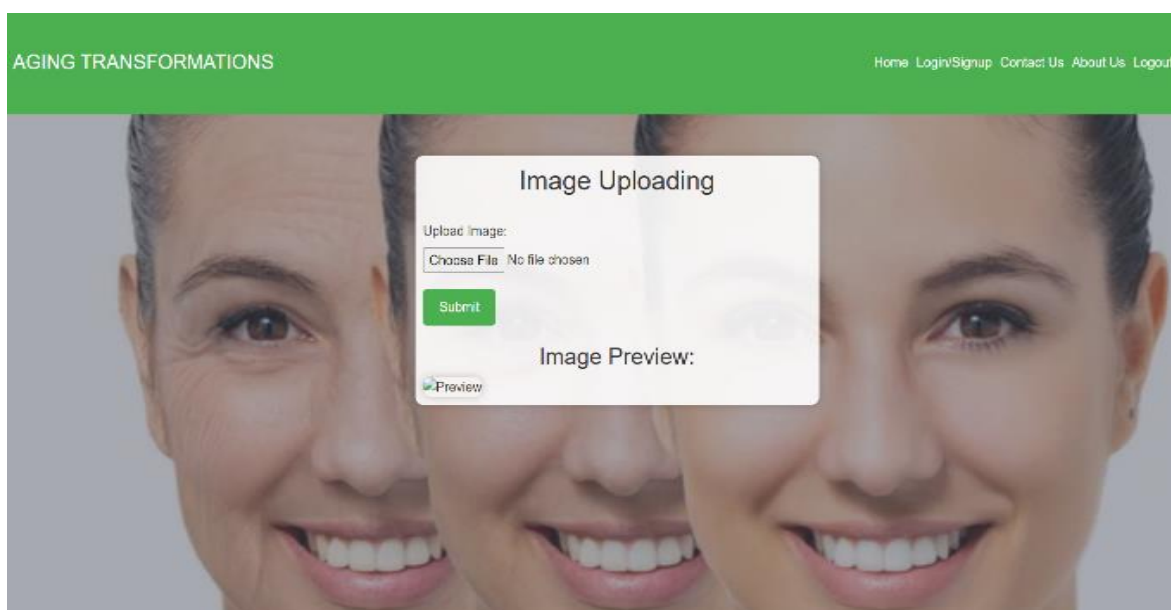


Figure 5. User will upload their single photograph.

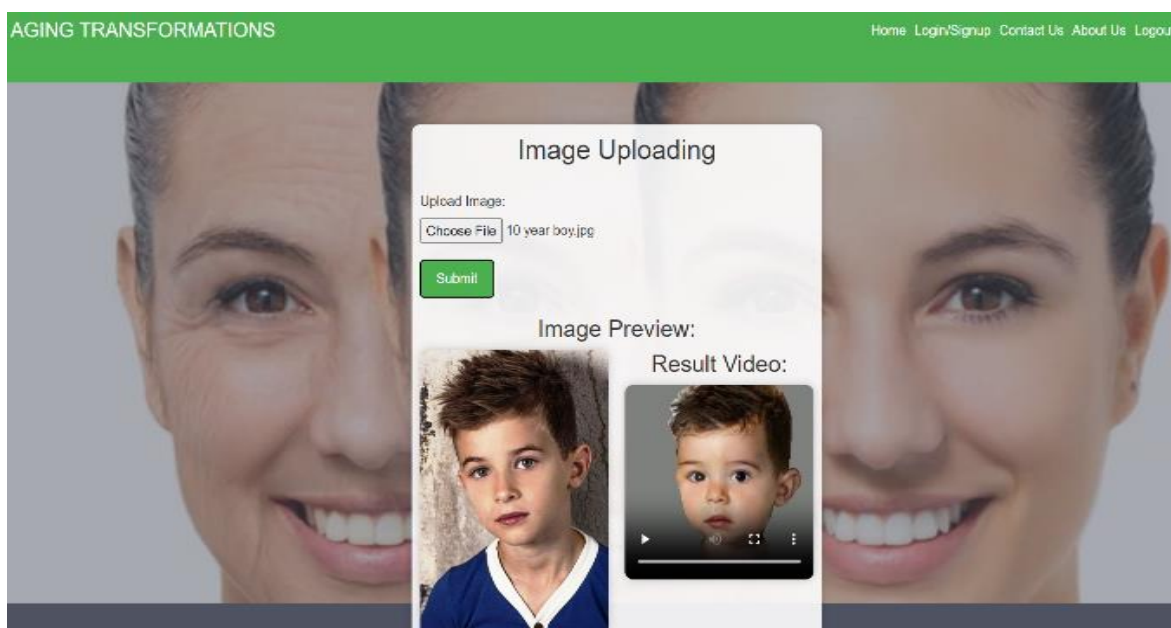


Figure 6. User can see the video clip of their age transformation (male).

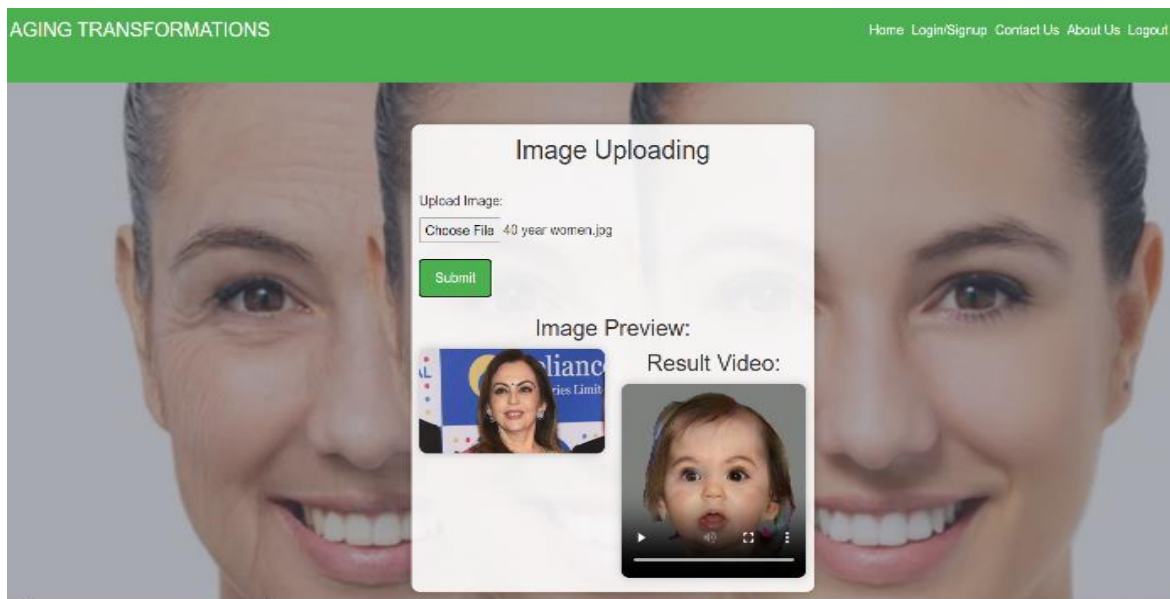


Figure 7. User can see the video clip of their age transformation (female).

CONCLUSION

The system generates trustworthy age transformations for anyone between the ages of 0 and 70 years. Our framework, in contrast to earlier methods, learns to alter both texture and shape as it ages. By using latent space interpolation, we may generate results for ages that were never observed during training because the suggested architecture and training plan accurately generalise age and also released a new facial dataset that the vision community can utilise for a variety of applications. The approach yields state-of-the-art results, as the experiments have shown.

Acknowledgement

We are deeply grateful to our project mentor, Dr. PVS Lakshmi Jagadamba, Associate Professor and Department Head, for her direction and unwavering support during our major college project. She gave us an attitude of comprehension and encouragement that enabled us to complete our project effectively. The head of the department of Computer Science and Engineering, Dr. PVS Lakshmi Jagadamba has our deepest gratitude for her insightful advice and unwavering encouragement, which tremendously helped us to successfully complete our major project.

REFERENCES

1. Antipov G, Baccouche M, Dugelay JL. Face aging with conditional generative adversarial networks. In 2017 IEEE international conference on image processing (ICIP). 2017 Sep 17; 2089–2093.
2. He X, Chen T, Kan MY, Chen X. Trirank: Review-aware explainable recommendation by modeling aspects. In Proceedings of the 24th ACM international on conference on information and knowledge management. 2015 Oct 17; 1661–1670.
3. Nhan Duong C, Gia Quach K, Luu K, Le N, Savvides M. Temporal non-volume preserving approach to facial age-progression and age-invariant face recognition. In Proceedings of the IEEE international conference on computer vision. 2017; 3735–3743.
4. Wang W, Yan Y, Cui Z, Feng J, Yan S, Sebe N. Recurrent face aging with hierarchical autoregressive memory. IEEE Trans Pattern Anal Mach Intell. 2018 Feb 7; 41(3): 654–68.
5. Wang W, Cui Z, Yan Y, Feng J, Yan S, Shu X, Sebe N. Recurrent face aging. In Proceedings of the IEEE conference on computer vision and pattern recognition. 2016; 2378–2386.
6. Tang J, Li Z, Lai H, Zhang L, Yan S. Personalized age progression with bi-level aging dictionary learning. IEEE Trans Pattern Anal Mach Intell. 2017 May 17; 40(4): 905–17.

7. Kemelmacher-Shlizerman I, Suwajanakorn S, Seitz SM. Illumination-aware age progression. In Proceedings of the IEEE conference on computer vision and pattern recognition. 2014; 3334–3341.
8. Bando Y, Kuratate T, Nishita T. A simple method for modeling wrinkles on human skin. In 10th IEEE Pacific Conference on Computer Graphics and Applications, 2002. Proceedings. 2002 Oct 9; 166–175.
9. Boissieux L, Kiss G, Thalmann NM, Kalra P. Simulation of skin aging and wrinkles with cosmetics insight. In Computer Animation and Simulation 2000: Proceedings of the Eurographics Workshop in Interlaken, Switzerland, August 21–22, 2000. Springer Vienna. 2000; 15–27.
10. Burt DM, Perrett DI. Perception of age in adult Caucasian male faces: computer graphic manipulation of shape and colour information. Proc R Soc Lond B: Biol Sci. 1995 Feb 22; 259(1355): 137–43.
11. Chen BC, Chen CS, Hsu WH. Cross-age reference coding for age-invariant face recognition and retrieval. In Computer Vision–ECCV 2014: 13th European Conference, Zurich, Switzerland, September 6–12, 2014, Proceedings, Part VI 13. Springer International Publishing. 2014; 768–783.
12. Chen LC. Rethinking atrous convolution for semantic image segmentation. arXiv preprint arXiv:1706.05587. 2017.
13. Choi Y, Choi M, Kim M, Ha JW, Kim S, Choo J. Stargan: Unified generative adversarial networks for multi-domain image-to-image translation. In Proceedings of the IEEE conference on computer vision and pattern recognition. 2018; 8789–8797.
14. Choi Y, Uh Y, Yoo J, Ha JW. Stargan v2: Diverse image synthesis for multiple domains. In Proceedings of the IEEE/CVF conference on computer vision and pattern recognition. 2020; 8188–8197.
15. Duong CN, Luu K, Quach KG, Bui TD. Longitudinal face aging in the wild-recent deep learning approaches. arXiv preprint arXiv:1802.08726. 2018 Feb 23.
16. Fu Y, Guo G, Huang TS. Age synthesis and estimation via faces: A survey. IEEE Trans Pattern Anal Mach Intell. 2010 Feb 5; 32(11): 1955–76.
17. He K, Zhang X, Ren S, Sun J. Deep residual learning for image recognition. In Proceedings of the IEEE conference on computer vision and pattern recognition. 2016; 770–778.
18. Huang X, Belongie S. Arbitrary style transfer in real-time with adaptive instance normalization. In Proceedings of the IEEE international conference on computer vision. 2017; 1501–1510.
19. Huang X, Liu MY, Belongie S, Kautz J. Multimodal unsupervised image-to-image translation. In Proceedings of the European conference on computer vision (ECCV). 2018; 172–189.
20. Isola P, Zhu JY, Zhou T, Efros AA. Image-to-image translation with conditional adversarial networks. In Proceedings of the IEEE conference on computer vision and pattern recognition. 2017; 1125–1134.