

Age Invariant Face Recognition using Lightweight Deep Convolutional Neural Networks

Rediet Tadesse

HiLCoE School of Computer Science and Technology
Software Engineering Programme
rediettadesse100@gmail.com

Fekade Getahun

Department of Computer Science, Addis Ababa University
fekade.getahun@aau.edu.et

Abstract

Age Invariant Facial Recognition is a field of study that is constantly evolving and is applied in a variety of real-world settings. The majority of age-invariant facial recognition works that make use of Convolutional Neural Networks rely on heavyweight Convolutional Neural Networks. However, training and inferring with heavyweight networks is an extremely resource-consuming task. Lightweight models have a greater advantage over the heavyweight ones, as they have less computational and memory (in terms of the number of parameters) demands, and also are proven to perform well in face recognition tasks. But no previous comparison has been made on which Lightweight Architecture is best suited for the Age Invariant Face Recognition task. Hence, this research aims at achieving three objective 1) compare the performance of 5 pre-trained lightweight Deep Convolutional Neural Network Architectures/models for the task of Age Invariant Face Recognition. 2) conducting set of experiments with different techniques before or during classification to improve the recognition accuracy and, 3) compare the performance of lightweight models pre-trained only on ImageNet for the task of AIFR.

After conducting all experiments for objectives 1 and 2, the highest accuracies that are achieved by each of the models are 80.75% for MobileNetV1, 87.7% for MobileNetV2, 84.88% for EfficientNet, 80.65% for ShuffleNet, and 78.43% for SqueezeNet. This shows that the MobileNetV2 (pre-trained on a face dataset) is the best-suited architecture/model for the task of Age Invariant Face Recognition. The results of the experiments also show that combining the original embedding's of the classifiers training set with embedding's of Generative Adversarial Networks/FaceApp augmented images leads to an increase in classification accuracy given that (1) A Support vector machine classifier with a "Radial Basis Function" kernel is used and (2) Each embedding is represented individually (the highest accuracy is achieved when it is used along with gender guidance). Moreover it is also shows that fine-tuning leads to an increase in accuracy. The experiments conducted to achieve objective 3 show that shows that lightweight models pre-trained only on ImageNet

are not suitable for the task of AIFR.

Keywords: Lightweight CNN, Age Invariant Face Recognition, Pre-trained model, FaceApp, GAN, Support Vector Machine, ImageNet.

1. Introduction

Face Recognition is a technology that has shown tremendous advances in the past few years. Despite its advance, efficient face recognition still faces major challenges [1] caused by illumination, pose, expression, and aging [2, 3, 4, 5]. This research focuses on the problem of age variation in face recognition. The effect of aging causes significant intra-class variations [6, 7]; and hence images of the same identity will have lower similarities. Age-invariant facial recognition (AIFR) is one of the techniques used to tackle this problem. Various solutions have been proposed for Age Invariant Face Recognition; generative approach, use of handcrafted/engineered features such as Scale Invariant Feature Transform (SIFT), Local Binary Pattern (LBP) and its variants, etc. [6, 8], and use of deep Convolutional Neural Network (CNN) based features [1, 3, 4, 6, 8].

Deep CNNs have powerful feature-extracting capabilities, with high accuracy, and are reported as adequate frameworks in image classification and recognition tasks [6]. Moreover, various heavyweight and lightweight deep CNN architectures, depending on the number of configurable parameters and computational complexity, have been designed for the task of image recognition. Deep learning is extremely resource consuming (processor, memory, energy, etc.) [9], whereas the lightweight CNNs have less computational and memory (in terms of the number of parameters) demands. This research focuses on using pre-trained lightweight architectures for the task of AIFR.

The problems investigated in this research are 1) A thorough comparison of different well-known lightweight deep CNN architectures for the task of AIFR has not been carried out, hence it is not known which lightweight architecture is best suited for the task of AIFR. 2) The complexity of lightweight models is less when compared to heavyweight models, and theoretically, that leads to a lower accuracy of the lightweight models [10] unless the lightweight model's architecture is cleverly designed to avoid the reduction in accuracy, thus it is always necessary to investigate on additional techniques that can be applied to increase the recognition accuracy on AIFR tasks. 3) Although using ImageNet pre-trained models for various kinds of recognition tasks is common nowadays, not enough investigation has been done on the use of ImageNet pre-trained lightweight deep CNN models for the task of Age Invariant Face Recognition (AIFR).

The above concerns lead to the following research questions:

- Which of the well-known lightweight CNN architectures is well suited for age-invariant face recognition in terms of accuracy and computational complexity?

- What techniques could be applied before or during the classification stage in order to improve the recognition accuracy when lightweight models are used?
- Which of the well-known lightweight deep CNN architectures pre-trained only on ImageNet is well suited for the task of AIFR?

2. Related Work

In this section, the main works related to AIFR that make use of CNNs are presented.

In [6], Yousaf *et al.* compared the performance of nine diverse ImageNet pre-trained CNNs fine-tuned on the FG-NET-AD dataset. Their output shows that employing the Resnet-18 model for feature extraction and classification is found to be the most robust, time-efficient, and computationally effective model. In [2] Khiyari and Wechsler proposed a set-based matching approach where the probe and gallery templates are treated as collections of images rather than singletons; extracted features are grouped as sets to form the biometric templates of different subjects. The distance between subjects is the similarity distance between their respective sets. Better performance is reported in generalization due to transfer learning, and local processing due to the combined use of CNN and robust similarity distances for set images rather than singletons. Wen, Li and Qiao in [11] proposes a novel deep face recognition framework to learn age-invariant deep face features through a carefully designed CNN model. Extensive experiments are conducted on several public domain face aging datasets (MORPH Album 2, FGNET, and CACD-VS) to demonstrate the effectiveness of the proposed model over the state-of-the-art.

Wang *et al.* in [7] proposes a novel approach named, Orthogonal Embedding CNNs, to learn age-invariant deep face features. Extensive experiments conducted on three public domain face aging datasets (MORPH Album 2, CACD-VS, and FGNET) have shown the effectiveness of the proposed approach.

In [4] Bianco proposes a new deep CNN architecture having a feature injection layer that fuses external features with the activation of the deepest deep CNN layers. Experimental results on the Large Age-Gap (LAG) dataset show that the proposed approach outperforms the face verification methods in the state of the art considered.

Bodhe, Kapse and Singh in [12] proposes a ScatterNet Inception Hybrid Network (SIHN) for AIFR, it consists of a hand-crafted ScatterNet front-end and an Inception ResNet-v1 based backend. Experimental results evaluated over 27,000 frames show that the proposed method achieve state-of-art performance on both video datasets as well as the other widely used datasets for age-invariant face datasets such as CACD and FG-NET.

In summary, it has seen that various techniques have been applied to address the

problem of AIFR. Some of the drawbacks observed are, most of the researches use heavyweight CNN's [2, 4, 6, 12] as feature extractors, despite the fact that these architectures are resource-consuming by nature, very little work has been done in the use of Lightweight CNN [6] models for the task of AIFR. Additionally, it is also noted that enough investigation of Lightweight models only pre-trained on ImageNet for the task of AIFR has not been performed. Moreover [6] performs the evaluation of ImageNet pre-trained CNN models on the FGNET dataset, but the evaluation strategy used is the 80% train and 20% test split set, and that would not be a fair comparison with other researches, as most researches evaluating on FGNET use the Leave one out testing strategy, some other researches that use this evaluation technique are [13], [14], [15], [16], [17], and [18]. Hence based on the drawbacks mentioned above, this research sets out to compare the performance of 5 pre-trained lightweight deep CNN models that are suited for resource-constrained devices, for the task of AIFR, it also aims to determine the best technique that could be applied before or during classification stage in order to improve the recognition accuracy. The research is limited to 5 models due to computational constraints, and the selection of these models is based on the number of parameters, computational complexity, availability of an ImageNet pre-trained model and its accuracy.

3. Proposed Approach

3.1. Proposed Approach in the Case of CASIA-Webface pre-trained models

i. Proposed architecture when embedding's of GAN/FaceApp augmented images are used to enhance the classifiers train set.

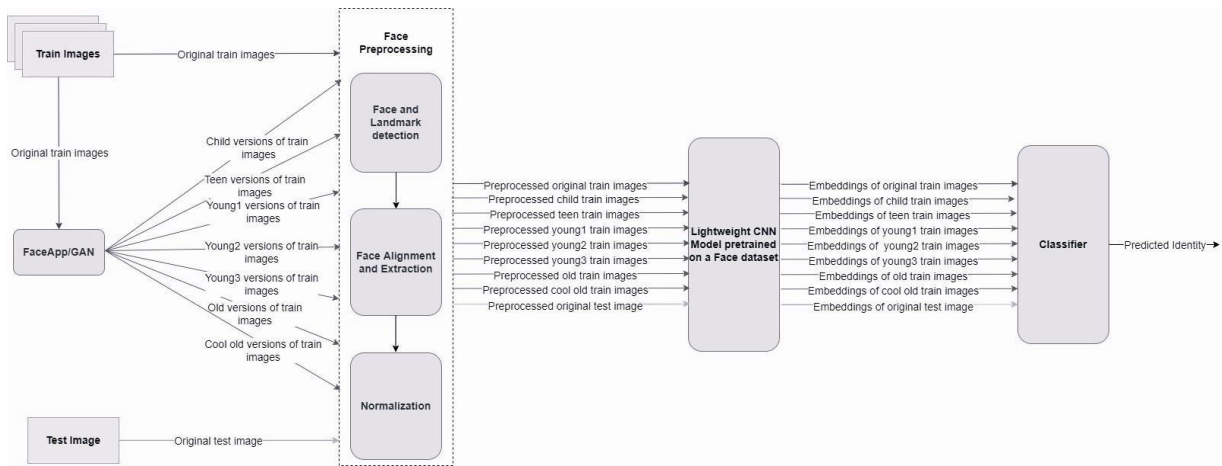


Figure 1. Proposed architecture when embedding's of GAN/FaceApp augmented images are used to enhance the classifiers train set

- ii. Proposed architecture when the CNN embedding's are concatenated with Principal Component Analysis (PCA) embedding's. (Experiments on embedding modifications before evaluation/classification is performed)

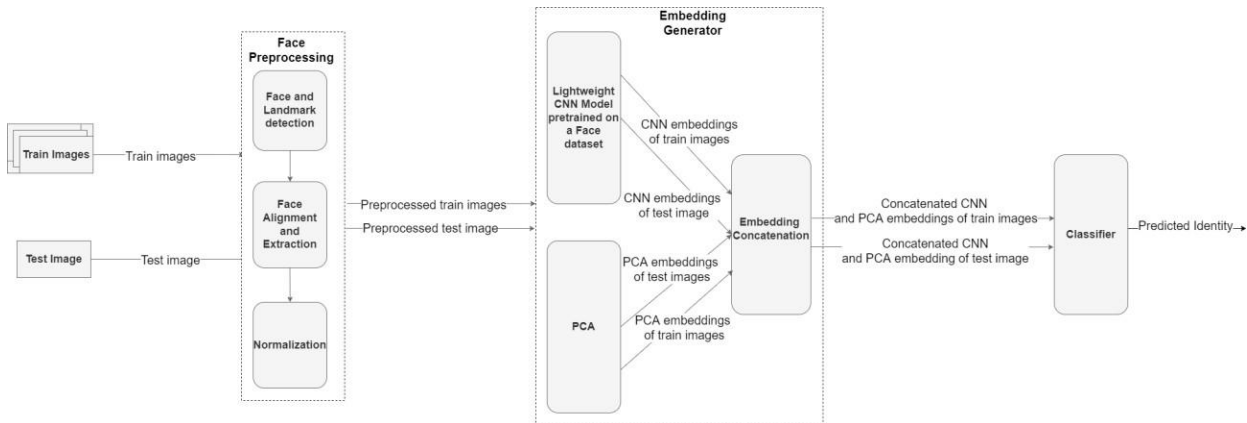
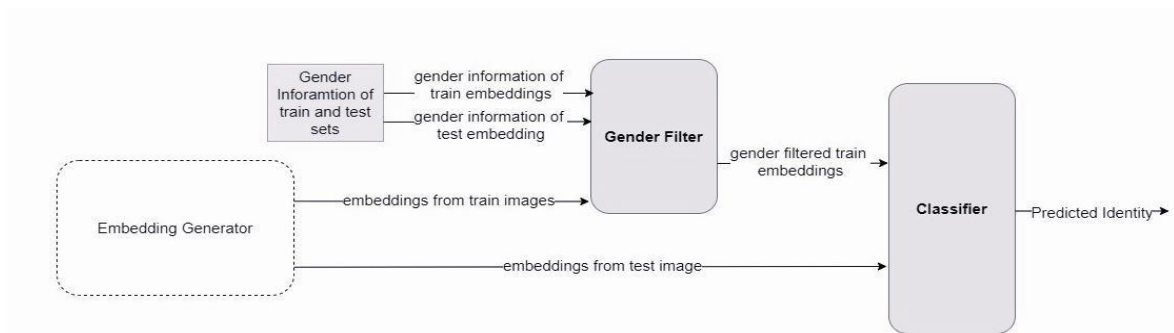


Figure 2. Proposed architecture when the CNN embedding's are concatenated with PCA embedding's

- iii. Proposed Architecture when gender guidance is used (Experiments on the evaluation process)

Gender guidance is introduced just before the classification stage, and it can be added to any of the above architectures, the diagram below shows the parts that will be added



- iv. Architecture for summing embedding's (Experiments on embedding representation before evaluation/classification is performed.)

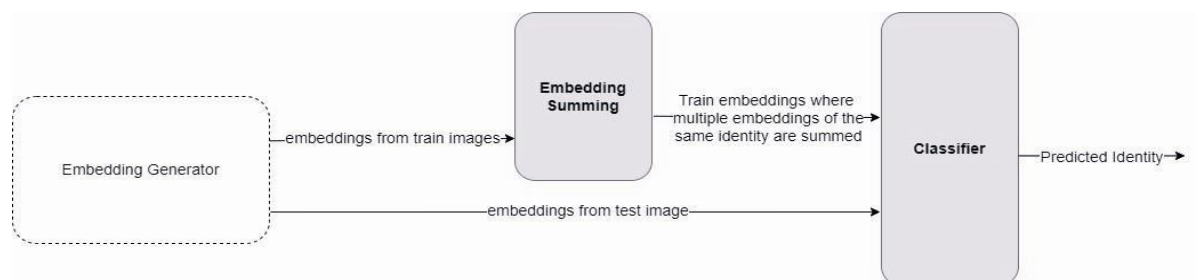


Figure 4. Architecture for summing embedding's

The embedding summer is introduced just before the classification stage, it sums multiple train embedding's of an individual. Any of the above architectures, including gender guidance, can incorporate it. This summing technique is observed in [19] and [20], but in this research it will be used in combination with the above architectures.

In order to perform the experiments, five well-known lightweight CNN models are selected; they are then prepared by removing their top layers and adding a few extra layers to the end of every model, which includes Global Depthwise Convolution (GDC) layers. These models are then pre-trained on a face dataset (CASIA-Webface) using the code and recommended hyper-parameters provided in [21] (with minor changes on embedding shape, output layer setup, and batch_size; 256, GDC, and 512 were chosen respectively). In addition, the selected loss functions are ArcFace, for MobileNet, MobileNetV2, EfficientNet, ShuffleNet, and Categorical Cross-entropy for SqueezeNet. Evaluation is then done using the LFW dataset using the evaluation script of Leondgarse [21]. The obtained LFW accuracies are 97.98% for MobileNet, 98.48% for MobileNetV2, 98.28% for EfficientNetB0, 96.28% for SqueezeNet-v1.1 and 97.52% for ShuffleNet.

Using the leave-one-out testing strategy with an embedding size of 256, the pre-trained models are then evaluated on the FGNET dataset for the task of age-invariant face recognition (the trainability of all layers is set to false during evaluation, except for the fine-tuning experiments). Moreover, the FGNET dataset is preprocessed (face alignment, extraction, and normalization) before performing the evaluation; the extracted faces used for the CASIA Webface models have a size of 112x112 before normalization is performed).

Various experiments/evaluations are performed to test the performance of each of the pre-trained models on the FGNET dataset. The performed experiments are listed below.

- Experiments on enhancing the classifiers train set with embedding's of Generative Adversarial Network (GAN)/FaceApp generated images.
 - Images generated by GANs are selected because they are capable of producing life like images [22] that realistically mimic the training set, and FaceApp is a popular photo-editing program that is renowned for producing incredibly life like changes to human faces in images. Many sources assert that the technology behind FaceApp is GAN's [23, 24, 25, 26].
 - In this experiment seven FaceApp aging filters (Child, Teen, Young1, Young2, Young3, Old, and Cool Old) are used to generate the images, and their embeddings are used to enhance the classifiers train set.
 - Experiments on embedding modifications before evaluation/classification is

performed.

- Experiments on embedding representation before evaluation/classification is performed.
- Experiments on the evaluation process (Gender Guidance).
- Experiments on the classifiers (K-Nearest Neighbors (KNN), Support Vector Machine (SVM), and Fully Connected Layers)
- Experiments on the models (Fine Tuning)

The above experiments are performed by introducing different modifications. The modifications introduced are shown in Figures 1, 2, 3 and 4.

The experiments listed above, excluding the fine-tuning experiment, are combined in different ways in order to know which modification strategy gives the best accuracy. A total of 37 experiments are performed for a single model, hence there are a total of 185 experiments for the five models. All the evaluation is done using the leave-one-out evaluation strategy, moreover, the embedding's used for experimentation are extracted from the second-to-last layer, which has a size of 256.

Once the best-performing model and strategy are selected, the model is further fine-tuned to see if that causes an increase in accuracy. This is done by setting the trainability of selected layers to true. To have a fair comparison with literature, the model that performed best in the non-gender- guided case is also fine-tuned.

3.2. Proposed Approach in the Case of ImageNet pre-trained models

ImageNet pre-trained models are downloaded for each of the 5 architectures, their accuracies are 70.4% for MobileNet [27], 71.3% for MobileNetV2 [27], 77.10% for EfficientNetB0 [27], 52.28% for ShuffleNet (Groups = 3) [28], and 57.5% for SqueezeNet-v1.1 [29] (Approximate value taken from the official SqueezeNet paper, as the exact accuracy is not stated).

The top layers of these models are removed and a Global average pooling (GAP) layer is added to the end of every model. After setting the trainability of all layers to false, the performance of these models is evaluated on the FGNET dataset.

Three of the experiments listed in section 3.1 are performed in order to test the performance of ImageNet-only pre-trained models on the FGNET dataset (for the task of AIFR). These are experiments on embedding representation before evaluation/classification is performed, experiments on the evaluation/classification process, and experiments on the classifiers. They are combined in different ways in order to know which modification strategy gives the best accuracy. Thirteen experiments are

performed for a single model, hence a total of 65 experiments are done for the 5 models. The model that performs best is then further fine-tuned. The preprocessed FGNET images used for the ImageNet pre-trained models have a final size of 224x224 before normalization is performed.

4. Discussion

4.1. Number of Parameters and Computational Complexity

Table 1. Total Number of Parameters and FLOPs of CASIA-Webface and ImageNet pre-trained models

Model	Total Number of Parameters		Computational Complexity (Floating Point Operations, FLOPs)	
	CASIA-Webface Pre-trained Models	ImageNet Pre-trained Models	CASIA-Webface Pre-trained Models	ImageNet Pre-trained Models
MobileNet	3,505,344	3,228,864	0.277G	1.15 G
MobileNetV2	2,612,288	2,257,984	0.166G	0.613 G
Efficient	4,403,875	4,049,571	0.221G	0.8 G
ShuffleNet	1,203,736	937,752	0.0782G	0.295 G
SqueezeNet-v1.1	875,072	722,496	0.121G	0.531 G

From the results shown in Table 1, we can observe that SqueezeNet-v1.1 is the lightest model in terms of the number of parameters; it is followed by ShuffleNet, MobileNetV2, MobileNet, and EfficientNet. From the same table, we can observe that the ShuffleNet model performs the least amount of computations. It is followed by SqueezeNet-v1.1, MobileNetV2, EfficientNet, and MobileNet. Therefore we can claim that ShuffleNet is the lightest in terms of computational complexity.

4.2. Accuracy Results

4.2.1. Accuracy Results for CASIA-Webface Pre-trained Models

The highest accuracies achieved for each of the models are shown in Figure 5. From Figure 5, we can observe that the MobileNetV2 architecture provides the maximum accuracy in two embedding cases: original embedding's combined with the embedding's of GAN/ FaceApp augmented images and original embedding's alone. Moreover, in both cases, the EfficientNetB0 architecture gives the second-highest accuracy. In the third case when the original embedding's are concatenated with the PCA embedding's, the SqueezeNet-v1.1 model gives the highest accuracy followed by the MobileNetV2 model.

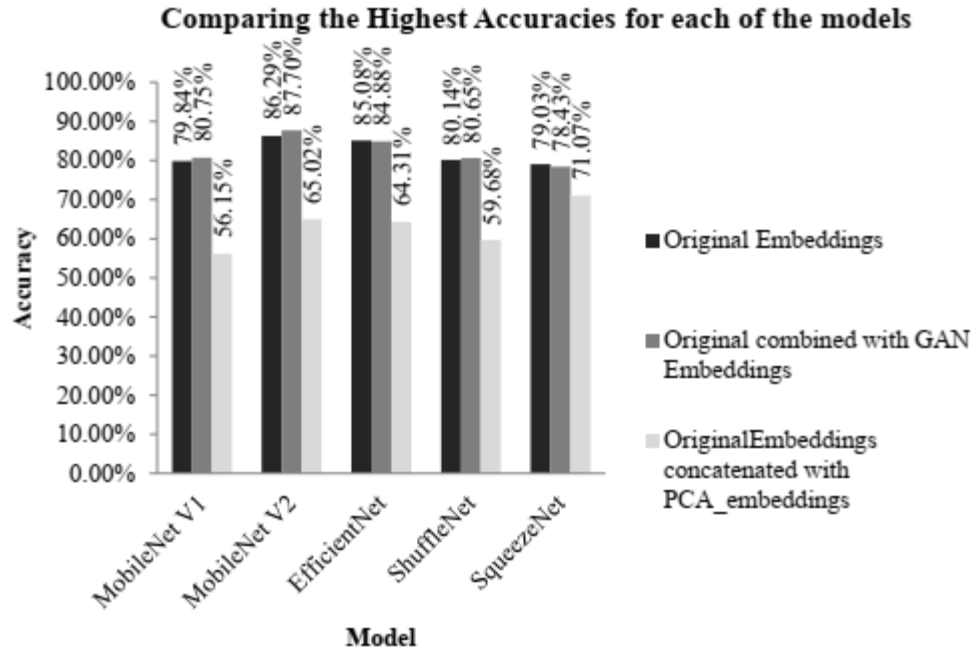


Figure 5. Comparing the highest accuracy in each of the models

Comparing all three cases, original embeddings only, original embeddings combined with GAN embeddings, and original embeddings concatenated with PCA embeddings, the highest accuracy, 87.7%, is achieved when the original embedding's are combined with the embedding's of GAN/ FaceApp augmented images, using the MobileNetV2 architecture. This accuracy is achieved using the following setup, (1) Combining the original embeddings of the classifiers train set with their corresponding GAN/FaceApp embeddings (embeddings of the corresponding GAN/FaceApp images), (2) Using an SVM classifier with an rbf kernel, (3) Each embedding is represented individually during classification/recognition (Embeddings of the same Identity are not summed) and (4) Gender guidance is introduced during the evaluation process

Other observations made are,

- In all models when the original embeddings are combined with the embeddings of GAN/FaceApp generated images the gender-guided version of the SVM classifier with an rbf kernel always has the best performance. But this is the case when embeddings are represented individually.
- In all models when only the original embeddings are used, the gender-guided version of the SVM classifier with a linear kernel always has the best performance. But this is the case when embedding's are represented individually. This is also true when the original embedding's are concatenated with the PCA embedding's.

- Using gender guidance during classification/recognition always leads to an increase in accuracy.
- Directly concatenating the original embedding's and the PCA embedding's leads to a significant decrease in accuracy compared to using the original embedding's only.
- Except in a few instances, expressing each embedding individually improves accuracy compared to summing all the embedding's of an identity.

Due to high computational requirements, the fine-tuning experiment is performed on the model that performed best using the original embedding setup, which has an accuracy of 86.29% (The setup used in this case is similar to the best-performing model with the only difference being that the SVM classifier with the linear kernel gave a higher accuracy instead of the rbf kernel). The accuracy before fine-tuning is 86.29%, and after fine-tuning this increased to 88.21%. The results from the fine-tuning experiments can be seen in Figure 6.

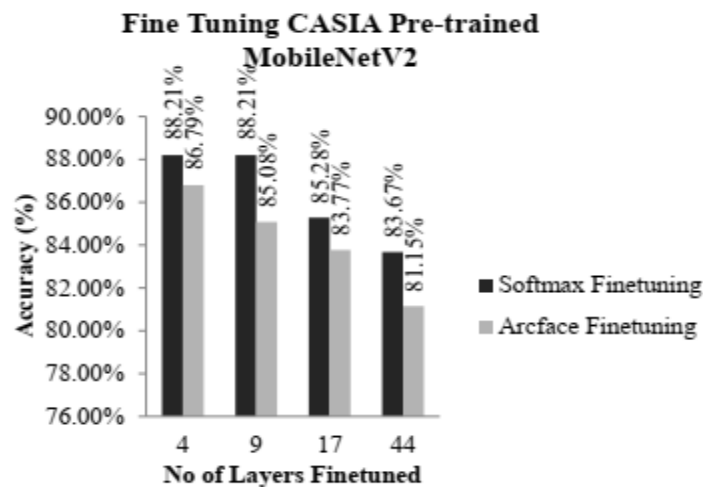


Figure 6. Accuracy results of fine-tuning the MobileNetV2 model

Moreover, from the Fine Tuning results it can be observed that the highest increase in accuracy is achieved when 4 or 9 numbers of layers are fine-tuned using categorical cross entropy as loss. We can also observe that, as the number of fine-tuned layers increases the accuracy starts to decrease, and using categorical cross-entropy as loss during fine-tuning leads to a higher accuracy. In relation to results reported on the FGNET dataset using the leave one out testing strategy, this accuracy of 88.21% shows good performance. The results reported in Literature are 93.20% by [13], 88.10% by [11], 86.50% by [14], 76.2% by [15], 69% by [16], 47.5% by [17], and 37.4% by [18]. For further comparison with literature, the non-gender guided version of the model that performed best using the original embedding setup is also fine-tuned. This is done so that we can know how well the model compares with the literature in cases it is not provided with gender information.

Originally it had an accuracy of 82.86%, after fine-tuning, the highest accuracy achieved is 86.09%, comparing it to the results of literature listed above, it can be seen that the accuracy achieved could be considered very good.

4.2.2. Accuracy Results for ImageNet Pre-trained Models

The accuracy achieved using ImageNet-only pre-trained models is very low; the highest accuracy achieved is 36.9% by the MobileNet model. The experimental setup used to obtain this accuracy is, (1) SVM classifier with Linear Kernel, (2) Use of gender guidance, and (3) Representing embedding's individually. Using the same setup, the model is further fine-tuned to see if an increase in accuracy is achieved, but the highest accuracy that could be achieved is 38.51%. This leads to the conclusion that models pre-trained only on ImageNet have poor performance for the task of Age Invariant Face Recognition.

5. Conclusion

This work focused on evaluating the performance of 5 pre-trained lightweight CNN models for the task of AIFR. The comparison was also done using different experimental setups in order to determine a setup/technique that could lead to higher classification accuracy. The results show that MobileNetV2 pre-trained with the CASIA-Webface dataset has the highest accuracy (87.7%). From the experimental setup that is used to achieve this accuracy, we can conclude that combining the original embedding's of the classifiers train set with the corresponding embedding's of GAN/FaceApp augmented Images, along with using an SVM classifier with rbf kernel, and representing the embedding of each individual separately leads to an increase in accuracy and the highest accuracy is achieved when gender guidance is used.

Due to high computational requirements, the fine-tuning experiment is performed on the model that performed best using the original embedding setup, MobileNetV2 with 86.29% accuracy. Upon fine-tuning 4 or 9 layers of this model using categorical cross entropy as loss, its accuracy increased from 86.29% to 88.21%. Other conclusions that can be made from this research are; models pre-trained only on ImageNet are not suitable for the task of AIFR, concatenating the original embedding's directly with PCA embedding's leads to significant decrease in accuracy, compared to the other classifiers SVM classifier with an rbf kernel has the best accuracy when original embedding's are combined with the embedding's of GAN/FaceApp augmented images, but this is the case when embedding's are represented individually and the best performance is seen by the gender guided version, compared to other classifiers SVM classifier with a linear kernel always has the best accuracy when only the original embedding's are used, but this is the case when embedding's are represented individually and the best performance is seen by the gender guided version, using gender guidance always leads to an increase in accuracy, and except in a few cases representing the embedding's individually leads to higher accuracy instead of summing them.

References

- [1] DivyanshuSinha, J. P. Pandey, and BhaveshChauhan, “A Deep Learning Approach for Age Invariant Face Recognition,” *Int. J. Pure Appl. Math.*, vol. 117, no. 21, pp. 371–389, 2017.
- [2] H. El Khiyari and H. Wechsler, “Age Invariant Face Recognition Using Convolutional Neural Networks and Set Distances,” *J. Inf. Secur.*, vol. 08, no. 03, pp. 174–185, 2017.
- [3] M. Nimbarte and K. Bhoyar, “Age invariant face recognition using convolutional neural network,” *Int. J. Electr. Comput. Eng.*, vol. 8, no. 4, pp. 2126–2138, 2018.
- [4] S. Bianco, “Large age-gap face verification by feature injection in deep networks,” pp. 1– 7, 2016.
- [5] M. Sajid *et al.*, “Demographic-Assisted Age-Invariant Face Recognition and Retrieval,” 2018.
- [6] A. Yousaf, M. J. Khan, M. J. Khan, and K. Khurshid, “A Robust and Efficient Convolutional Deep Learning Framework for Age Invariant Face Recognition,” 2020.
- [7] Y. Wang *et al.*, “Orthogonal Deep Features Decomposition for Age-Invariant Face Recognition,” vol. 1, pp. 1–13, 2018.
- [8] A. Dhamija and R. . Dubey, “Analysis on Age Invariance Face Recognition Study and Effects of Intrinsic and Extrinsic Factors on Skin Ageing,” *Int. J. Comput. Appl.*, vol. 182, no. 43, pp. 1–9, 2019.
- [9] C. Chen *et al.*, “Deep Learning on Computational-Resource-Limited Platforms: A Survey,” *Hindawi Mob. Inf. Syst.*, vol. 2020, 2020.
- [10] Z. Yassir, “(3/12) MobileNets: Depthwise Separable Convolutions (Part 2),” *Youtube*, 2021. [Online]. Available: https://www.youtube.com/watch?v=nc3_5bhE8f0. [Accessed: 24-Oct-2022].
- [11] Y. Wen, Z. Li, and Y. Qiao, “Latent factor guided convolutional neural networks for age- invariant face recognition,” *Proc. IEEE Comput. Soc. Conf. Comput. Vis. Pattern Recognit.*, 2016.
- [12] S. Bodhe, P. Kapse, and A. Singh, “Real-time age-invariant face recognition in videos using the scatternet inception hybrid network (SIHN),” *Proc. - 2019 Int. Conf. Comput. Vis. Work. ICCVW 2019*, 2019.
- [13] J. Zhao *et al.*, *Look Across Elapse: Disentangled Representation Learning and Photorealistic Cross-Age Face Synthesis for Age-Invariant Face Recognition*. 2018.
- [14] C. Xu, Q. Liu, and M. Ye, “Age invariant face recognition and retrieval by coupled auto- encoder networks,” *Neurocomputing*, vol. 222, pp. 62–71, 2017.
- [15] D. Gong, Z. Li, D. Tao, J. Liu, and X. Li, “A maximum entropy feature descriptor for age invariant face recognition,” in *2015 IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, 2015, pp. 5289–5297.
- [16] D. Gong, Z. Li, D. Lin, J. Liu, and X. Tang, “Hidden Factor Analysis for Age Invariant Face Recognition,” in *2013 IEEE International Conference on*

- Computer Vision*, 2013, pp. 2872–2879.
- [17] Z. Li, U. Park, and A. K. Jain, “A Discriminative Model for Age Invariant Face Recognition,” *IEEE Trans. Inf. Forensics Secur.*, vol. 6, no. 3, pp. 1028–1037, 2011.
 - [18] U. Park, Y. Tong, and A. K. Jain, “Age-Invariant Face Recognition,” *IEEE Trans. Pattern Anal. Mach. Intell.*, vol. 32, no. 5, pp. 947–954, 2010.
 - [19] Leondgarse, “Keras Insightface,” *Github*, 2022. [Online]. Available: https://github.com/leondgarse/Keras_insightface/blob/master/eval_folder.py. [Accessed: 18-Jul-2022].
 - [20] Q. Meng, S. Zhao, Z. Huang, and F. Zhou, “MagFace: A universal representation for face recognition and quality assessment,” *Proc. IEEE Comput. Soc. Conf. Comput. Vis. Pattern Recognit.*, pp. 14220–14229, 2021.
 - [21] Leondgarse, “Keras Insightface,” *Github*, 2022. [Online]. Available: https://github.com/leondgarse/Keras_insightface. [Accessed: 18-Jul-2022].
 - [22] Coursera, “Generative Adversarial Networks (GANs) Specialization,” *Coursera*. [Online]. Available: <https://www.coursera.org/specializations/generative-adversarial-networks-gans>. [Accessed: 18-Jul-2022].
 - [23] D. Chakraborty, “IN DEPTH OF Faceapp,” *Medium*, 16-Apr-2020. [Online]. Available: <https://medium.com/analytics-vidhya/in-depth-of-faceapp-a08be9fe86f6>. [Accessed: 18-Jul-2022].
 - [24] V. Thakur, “How Faceapp works?,” *ScienceABC*, 19-Jan-2022. [Online]. Available: <https://www.scienceabc.com/innovation/how-does-faceapp-work.html>. [Accessed: 18-Jul-2022].
 - [25] ERA-IITK, “FaceApp,” *Medium*, 12-Aug-2019. [Online]. Available: <https://medium.com/@eraiitk/faceapp-81b0bb6ad0e1>. [Accessed: 18-Jul-2022].
 - [26] S. Das, “The AI Behind FaceApp,” *Analytics India Magazine*, 29-Jun-2020. [Online]. Available: <https://analyticsindiamag.com/the-ai-behind-faceapp/>. [Accessed: 18-Jul-2022].
 - [27] Keras, “Keras Applications.” [Online]. Available: <https://keras.io/api/applications/>. [Accessed: 18-Jul-2022].
 - [28] T. Scheck, “Keras ShuffleNet,” *Github*, 16-Oct-2020. [Online]. Available: <https://github.com/scheckmedia/keras-shufflenet>. [Accessed: 18-Jul-2022].
 - [29] R. C. Malli, “keras-squeezenet,” *Github*, 19-Feb-2019. [Online]. Available: <https://github.com/rcmalli/keras-squeezenet>. [Accessed: 18-Jul-2022].