

Deep Learning for Facial Forgery Detection

Performance Evaluation of DenseNet201, InceptionV3 and ConvNeXt

¹Akshatha G, ²Kempanna M, ³Ashoka S B and ⁴Job Prasanth Kumar Chinta Kunta

¹Department of Computer Science and Engineering, BMS College of Engineering, Bangalore, Visvesveraya Technological University, Belagavi, Karnataka, India.

²Department of Artificial Intelligence and Machine Learning, Bangalore Institute of Technology, Bangalore, Visvesveraya Technological University, Belagavi, Karnataka, India.

³Department of Computer Science, Maharani Cluster University, Bangalore, Karnataka, India.

⁴Lead Solution Data Architect, The Automobile Association, London, United Kingdom.

¹akshu965@gmail.com, ²kempsindia@gmail.com, ³dr.ashoksbcs@gmail.com, ⁴job.cintakunta@gmail.com

Correspondence should be addressed to Akshatha G: akshu965@gmail.com

Article Info

Journal of Machine and Computing (<https://anapub.co.ke/journals/jmc/jmc.html>)

Doi: <https://doi.org/10.53759/7669/jmc202606005>

Received 23 June 2025; Revised from 30 August 2025; Accepted 02 October 2025.

Available online 14 October 2025.

©2026 The Authors. Published by AnaPub Publications.

This is an open access article under the CC BY-NC-ND license. (<https://creativecommons.org/licenses/by-nc-nd/4.0/>)

Abstract – The recent spread of AI-generated face forgery is one of the greatest threats to visual media credibility. The Proposed work compares three deep transfer learning models DenseNet201, InceptionV3, and ConvNeXt, in detecting manipulated facial images. An 8,000 real and fake facial image dataset was used to train and benchmark models under consistent experimental condition. ConvNeXt achieved the best classification accuracy of 91.25 % which is much higher than that of DenseNet201 (75.12 %) and InceptionV3 (68.38 %). In addition, ConvNeXt had better trade-off between true positive and false positive rates, which means better generalization and resistance to overfitting. These findings prove the applicability of ConvNeXt in the robust detection of facial forgery and highlight potential application in the practical implementation of the facial authenticity determination. Future research will investigate ensemble methods and real-time inference.

Keywords – Facial Fake Images, Transfer Learning (TL), DenseNet201, InceptionV3, Facial Forgery Detection, ConvNeXt.

I. INTRODUCTION

The rise in deep learning tools like GANs, many fake facial images are now a challenge for forensics experts. This type of forged media called deepfakes can cause problems in areas such as spreading political lies, stealing someone's identity or committing fraud and social manipulation. For this reason, computer vision and artificial intelligence experts now pay special attention to accurate and reliable facial image detection. While several solutions have been developed to identify deepfakes, they often fail to work well under different conditions or in realistic situations. For instance, EfficientNet and dual-stream architecture has demonstrated improvements in identifying situations where the face does not match the context [1][2]. At the same time, hashing and texture filtering [3][4] are among the feature-based methods that target the special statistical characteristics of synthetic images. However, they face difficulties when dealing with new kinds of manipulations or when they are used in real time.

At the same time, newer approaches in face manipulation detection use attention, autoencoders and saliency-based methods [5][6]. Even so, most of these methods require a large amount of data to train or do not maintain their accuracy when things are noisy. Additionally, using semi-supervised and lightweight models is not widely explored for this purpose. ConvNeXt recently improved convolutional architecture and now outperforms other methods in many vision tasks. [7][8] have applied ConvNeXt in a semi-supervised framework with consistency regularization and [9] present a fast version of the model, ConvFaceNeXt, for face recognition. Researchers have proven that when attention mechanisms and architectural changes are applied to ConvNeXt, both fine-grained classification and anti-spoofing detection are enhanced. It seems that ConvNeXt is capable of effective facial deception detection when set up correctly. The study aims to assess and contrast the abilities of DenseNet201, InceptionV3 and ConvNeXt at recognizing both genuine and fake faces. Using a dataset that is publicly accessible, we tune each model within the same setup, study their performance by considering accuracy, precision, recall and F1-score and review their confusion matrices and how they learn.

The primary contributions of this paper are as follows:

- We use DenseNet201, InceptionV3 and ConvNeXt for this task by employing transfer learning.
- We analyze and compare the behavior, advantages and weaknesses of the models under equal test circumstances.
- We discuss matters involved in deploying these models, for example, overfitting, how well they can be used for different purposes and whether they can benefit from being trained together in groups.



Fig 1. Sample of Real and Fake Images.

This research seeks to design and deploy a facial image deception detection system by using transfer learning [10] and relying on the DenseNet201, InceptionV3 and ConvNeXt architectures. The purpose of this study is to teach and examine the models using a set of balanced images, so they can recognize what makes a facial photo real or fake. Evaluating the three models based on accuracy, precision, recall and F1-score is the main aim of the research. To determine the best method, the researchers analyze and record these metrics and consider if the approach is practical for real-world settings **Fig 1** shows Sample of Real and Fake Images.

The remainder of the paper is organized as follows: Section 2 presents the literature review; Section 3 presents the proposed methodology. Section 4 provides an overview of the results. Section 5 contains the work's conclusion.

II. LITERATURE REVIEW

In recent times, deep learning has improved forgery, deepfake and attack detection in images by applying both old CNN models and new ConvNeXt ones.

Deepfake Detection Techniques

They suggest a way of detecting deepfake text that is becoming increasingly common on social media. Using a combination of CNN and FastText embeddings, they managed to achieve 93 % accuracy in telling apart human-written and bot-written tweets. Thanks to this method, we understand how vital it is to fight misinformation on the internet and learn how to protect democracy. In a similar vein, [11] mention that it is still difficult to reconstruct images from deepfake faces, even with the EfficientNet model. They use a combination of context and face recognition to identify cases of identity manipulation and face swapping and the approach works well for unseen types of attacks. [12] compare several image forgery detectors against attacks performed by GAN based image [13] to image translation, while [14] make use of global texture to differentiate between real and GAN generated faces. In general, they point out the successes and the challenges in finding deepfakes.

Advancements in Fake Image Detection

Measures to address the problem of fake images in images are evolving by using new ways to display features and combinations of different techniques. LBP Net is introduced by [15] and makes use of local binary patterns to detect dissimilarities in texture between real and fake faces. A quantum inspired evolutionary feature selector and pre trained CNN are combined by [16] to improve the process of detecting fake faces. [17] merge the two types of feature extractors in their CNN model, making it robust even when working with video frames that are compressed. [18] suggest a resilient hashing algorithm to JPEG compression, while [19] rely on shallow CNNs applied to residual waveforms in multiple colors to detect GAN images with good accuracy.

Face Manipulation Detection

Advances in finding manipulated facial content are thanks to deep learning and CNNs. [20] use both error analysis and normalization for face image levels to train different pre trained models used to detect forged faces.

III. RESEARCH METHODOLOGY

The **Fig 2** demonstrates how to compare DenseNet201, ConvNeXt and InceptionV3 models for fake face detection. Data preparation involves a "Fake Face Dataset" like CIPL Lab & GreatGameDota present in Kaggle, followed by analysis, preparing the data and dividing it. Then, the models are trained, evaluated and compared based on accuracy, precision, recall, F1-score and the confusion matrix.

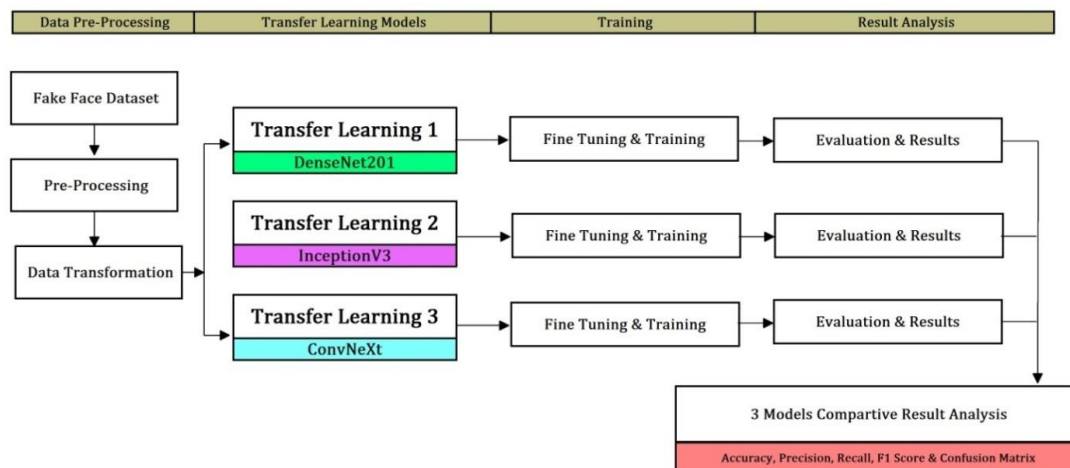


Fig 2. Proposed System Architecture.

Image Preprocessing

Getting your images ready for use in machine learning depends greatly on image pre-processing. It applies several techniques to enhance how suitable and easy-to-see images are for study.

Many times, adjusting image dimensions to a standard size by resizing is used in machine learning. It becomes important when images are of different sizes or when the model is designed for inputs of a specific size.

For training and evaluation, the Real and Fake Face Detection depends on two key datasets. The CIPL Lab on Kaggle gives us access to 1,081 actual face images and 960 fake face images. An additional 650 fake face images come from the second dataset, FaceForensics by GreatGameDota. Techniques that help with data imbalance by enhancing datasets were applied, resulting in 4,000 sample images of both real and fake faces. As a result of augmenting the data, it is now easier to train models accurately between original and altered facial pictures.

Scaling Image Values: Normalization is necessary to ensure that the pixel values in an entire image are in the same range. This helps every pixel in all images have the same value and makes it easier for the model to understand what's in the images. If the pixel values are not close to zero or there is a wide range of values in the images, normalization becomes necessary.

Split Test Train

This study uses a 8,000 image samples, and the same number of real and fake specimens. Leveraging the conventional 80:20 training-to-testing ratio promoted in past research, the current work uses a more precise 80:10:10 training-validation-testing split, which uses 80%, 10%, and 10% images respectively to training, validate, and test. This is an attempt to provide a more detailed measure of model performance and it provides an opportunity to perform a rigorous test of generalization by separating the pattern captured during training and the ability of the model to remain accurate when applied to new data.

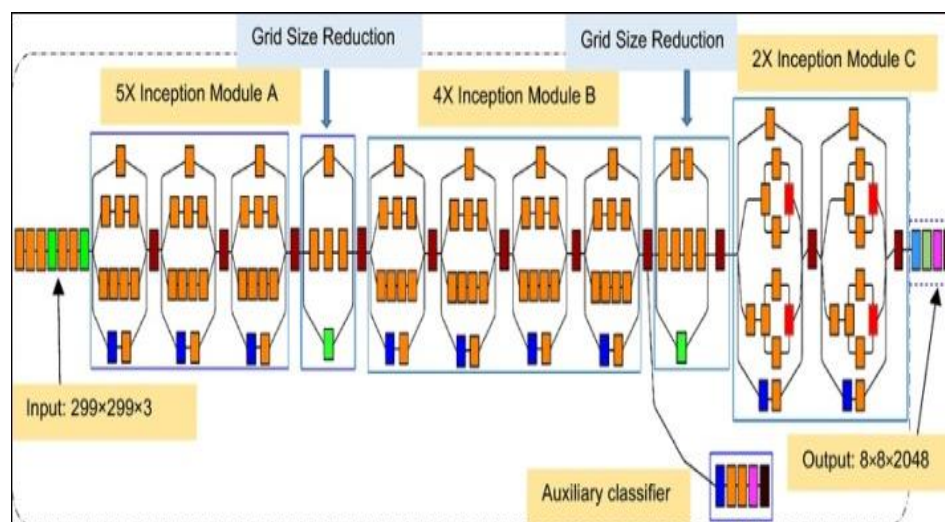


Fig 3. The Inception V3 Architecture.

Modeling

The method relies on deep learning (DL). The data will be separated into a test set and a training set. The training set will be used to train the model. Training begins once the model has been built. By using deep learning, we can develop classifiers that tell apart actual from fake facial photos.

We plan to use the dataset to develop three models: DenseNet201, InceptionV3 and ConvNeXt. At first, machine learning algorithms will be used to analyze all models and examine their results based on accuracy, precision and F1 score.

InceptionV3

The ImageNet datasets have allowed Inception v3 to show a remarkable 78.1% accuracy rate. A number of academics have spent years studying and testing ideas to develop this method. Inception v3 includes different types of elements, both symmetrical and asymmetric such as sequences, dropouts, fully connected layers, average and maximum pooling and convolution layers, in its design. All input data in the system is handled using a key method called batch normalization. The softmax function allows you to calculate losses which improves the accuracy of the model's forecasts.

DenseNet201

One CNN with 201 layers is known as DenseNet-201. This network is available with pretraining using ImageNet which holds a massive collection of training images. Having 1,000 different categories such as pens, keyboards, laptops and animals, this trained network is able to detect photographs **Fig 3** shows The Inception V3 Architecture. Thanks to its training, the network can now pick out small details from various pictures. DenseNet-201 is designed to process images with 224 by 224-pixel resolution. After getting information from other layers, each layer creates its own feature maps and passes them forward. Every layer is able to take the earlier information and use it to produce new findings, thanks to the technique of concatenation. Every layer uses knowledge from the layers above to help the network find more and more complex information as data moves forward.

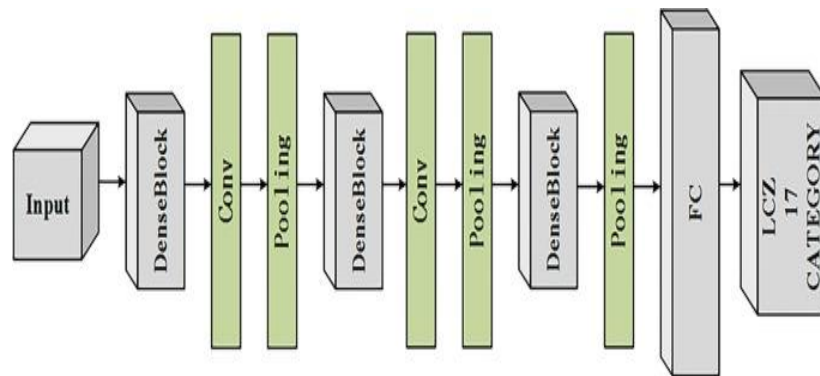


Fig 4. The DenseNet201 Architecture.

ConvNeXt

ConvNeXt connects the benefits of old-style convolutional neural networks with the new ideas from Vision Transformers, dividing its structure into stages that progressively reduce the size and enhance the details in feature maps. Each stage is designed with ConvNeXt blocks containing large kernel depth wise convolutions for reliable feature extraction which are then followed by layer normalization and a GELU-activated MLP that alternates the expansion and reduction of channel dimensions. Residual connections help the model perform at a high level with reduced computational cost and better gradient flow and stability **Fig 4** shows The DenseNet201 Architecture.

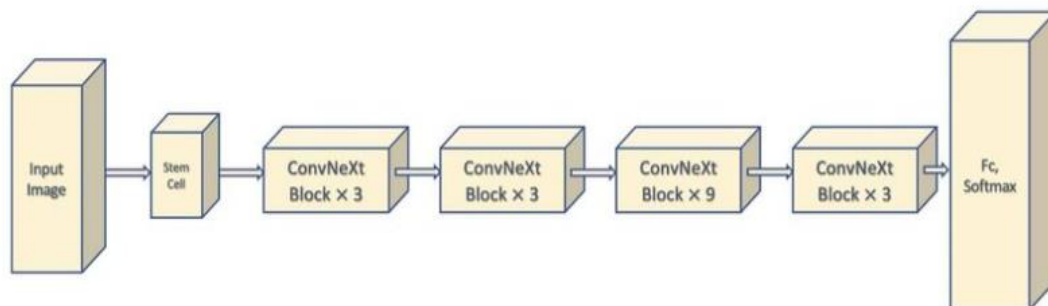


Fig 5. ConvNeXt Architecture.

Mathematical Formulations and Evaluation Metrics

This section adds to the previous modeling techniques which emphasizes by discussing the mathematical equations used to train and measure our DenseNet201, InceptionV3 and ConvNeXt deep learning models for facial forgery detection. The formulas and measurements given here are in line with the approaches used in the experiments described in this study **Fig 5** shows ConvNeXt Architecture.

Loss Function

It is explained that the binary cross-entropy loss is used to train the models which makes them suitable for telling apart real and fake facial images. The loss function can be written as:

$$\mathcal{L} = -[y \log(\hat{y}) + (1 - y) \log(1 - \hat{y})] \quad (1)$$

y represents the true label (1 for fake, 0 for real), $\hat{y}$ is the probability predicted by the model.

This is similar to the strategy used in robust image analysis tasks shown in our paper for instance.

Softmax Function (As Used in InceptionV3)

In the InceptionV3 model uses a softmax function to convert logits into a normalized distribution of probabilities. The softmax function is created with the following equation:

$$\hat{y}_i = \frac{e^{z_i}}{\sum_{j=1}^C e^{z_j}} \quad (2)$$

z_i is the logit (the unnormalized score) for class i , C is the total number of classes in our binary setting, $C=2$. This approach to transforming probabilities is in line with the details given about InceptionV3 in the document.

Performance Evaluation Metrics

To assess the performance of our models, we employ the following metric:

$$Accuracy = \frac{TP+TN}{TP+TN+FP+FN} \quad (3)$$

$$Precision = \frac{TP}{TP+FP} \quad (4)$$

$$Recall = \frac{TP}{TP+FN} \quad (5)$$

$$F1 = 2 \times \frac{Precision \cdot Recall}{Precision+Recall} \quad (6)$$

TP, TN, FP and FN stand for the true positives, true negatives, false positives and false negatives.

These metrics above help us understand how our model succeeds in identifying whether an image has been altered.

This information clearly outlines what the model can and cannot do in terms of spotting fake content.

Residual Connections (As Utilized in ConvNeXt)

Considering the idea of residual learning put forward in ConvNeXt, our network makes use of residual connections, as described in this research.

$$y = \mathcal{F}(x, \{W_i\}) + x \quad (7)$$

x is the input to the residual block, $\mathcal{F}(x, \{W_i\})$ represents the series of operations (e.g., convolution, normalization, and activation) applied to x . $\{W_i\}$ denotes the learnable parameters within the block.

IV. IMPLEMENTATION & RESULTS ANALYSIS

Here, DenseNet201, InceptionV3 and ConvNeXt were tested on their abilities to detect fake facial images by reviewing the classification reports and confusion matrices. They help point out the advantages and disadvantages of each model, giving essential information on how useful they could be when detecting fake faces.

InceptionV3 – Model

InceptionV3 architecture, achieved an accuracy of 68.38 % that is in line with high inter-class stability. However, the difference between training and validation accuracy and the apparent validation loss raise indicate the existence of

overfitting and the associated loss of trustworthiness of the image classification with respect to Fake class **Fig 6** shows Classification Report of InceptionV3 Model.

Classification Report for InceptionV3

	precision	recall	f1-score	support
Real	0.67	0.74	0.70	400
Fake	0.71	0.63	0.66	400
accuracy			0.68	800
macro avg	0.69	0.68	0.68	800
weighted avg	0.69	0.68	0.68	800

Fig 6. Classification Report of InceptionV3 Model.

Confusion Matrix for InceptionV3

		Confusion Matrix	
		Real	Fake
		Real 296	104
	Fake	149	251

Fig 7. Confusion Matrix for InceptionV3.

The confusion matrix displayed in **Fig 7** shows that, of the 400 Fake images, 251 were correctly classified but 149 of them were assigned to the Real category, resulting in 149 false negatives. In contrast, 400 Real images, 296 were correctly identified and 104 were incorrectly classified as Fake. The difference between the misclassification of the Fake and Real images indicates the lower reliability of the model, especially when it comes to the recognition of the fake instances, which implies that the performance in this category is lower.

DenseNet201 - Model

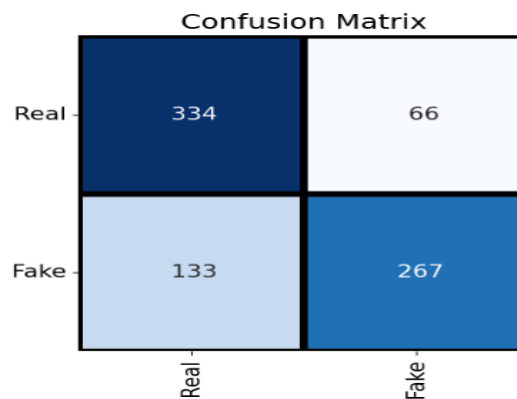
The experimental research with the use of the DenseNet201 architecture has reached the classification accuracy of 75.12%. The results of using Real images were overall stable, with the model performing less in generalizing Fake samples. There was a slight overfitting tendency beyond epoch 30.

Classification Report for DenseNet201

	precision	recall	f1-score	support
Real	0.72	0.83	0.77	400
Fake	0.80	0.67	0.73	400
accuracy			0.75	800
macro avg	0.76	0.75	0.75	800
weighted avg	0.76	0.75	0.75	800

Fig 8. Classification Report for DenseNet201.

The DenseNet model achieved overall accuracy of 0.75 on a test set of totals 800 samples with 400 Real and Fake samples each. As illustrated in **Fig 8** Fake class achieved a of precision 0.80, the recall 0.67 and the F1 score 0.73. The precision of the class Real was 0.72, recall 0.83, and F1 score was 0.77. The macro and weighted metrics averaged out over both classes were 0.76 and 0.75 respectively, which means that the classification performance was fair but slightly unbalanced.

Confusion Matrix for DenseNet201**Fig 9.** Confusion Matrix for DenseNet201.

DenseNet model has identified 267 of 400 Fake and mislabeled 133 of them as Real. In the Real class, 334 samples have been correctly classified and 66 were mislabeled as Fake. These results suggest that the model shows a definite asymmetry favoring false negatives (the number of false negatives is 133 whereas the number of false positives is 66) as illustrated in **Fig 9**. This disproportion emphasizes that the model is more accurate in differentiating Real samples than efficiency in identifying Fake items.

ConvNeXt – Model

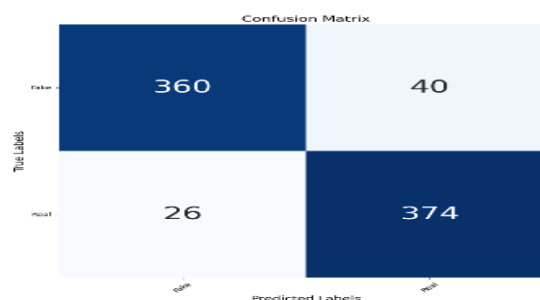
ConvNeXt architecture was tested thoroughly on its ability to distinguish between Fake and Real images and the results were identified as having sustained performance and minimal error, with a model accuracy of 91.75%. The model was able to exhibit strong and balanced generalization both in the training and validation phase.

Classification Report of ConvNeXt

Classification Report:		precision	recall	f1-score	support
Fake		0.93	0.90	0.92	400
Real		0.90	0.94	0.92	400
accuracy				0.92	800
macro avg		0.92	0.92	0.92	800
weighted avg		0.92	0.92	0.92	800

Fig 10. Classification Report of ConvNeXt.

The Classification Report of ConvNeXt architecture as illustrated in **Fig 10** proved to have strong discriminatory ability when used on the task of distinguishing Fake and Real images. In the case of Fake, the model had a precision of 0.93, that is, 93 % of the instances that the model predicted to be Fake were actually fake; the recall was 0.90, that is, 90 % of the actual Fake cases were correctly identified. The F1-score was 0.92, which means balanced performance. The Real class gave a precision of 0.90, a recall of 0.94, and F1-score of 0.92. The model was evaluated, with a total accuracy of 0.92. The macro and weighted averages of precision, recall, and F1-score were equal to 0.92, which emphasizes the stable and fair performance on both classes.

Confusion Matrix of ConvNeXt**Fig 11.** Confusion Matrix of ConvNeXt.

The ConvNeXt architecture's confusion matrix has demonstrated that out of 400 Fake samples, 360 of them were correctly identified, and 40 of them were misclassified as Real. Among the 400 real cases, 374 were classified correctly and only 26 were classified as Fake. Such imbalance between misclassifications is depicted in **Fig 11** where both the false positive and the false negative rates are lower, along with a balanced accuracy of 0.92. All these findings support the reliability of the model in distinguishing between Fake and Real.

Accuracy & Loss Plot of ConvNeXt

Accuracy and loss plots of the training and validation of the ConvNeXt model illustrated in **Fig 12** represent 30 training epochs each. Within the accuracy plot, the training accuracy begins at 0.7380 in epoch 1 and quickly goes up to 0.9890 within epoch 6. It is always high, as it is at epoch 10 (0.9960) and it maintains such a high quality (almost perfect) up to epoch 30 (approximately 0.9980). The accuracy of validation starts at 0.7910 and reaches the highest point of 0.9210 at epoch 6. It then oscillates a little bit between 0.9050 and 0.9180 up to epoch 30. This trend implies that as the model keeps training on training data, its performance on unseen data becomes stagnant at an early stage.

The plot of loss shows that the training loss (presented in blue) starts at 0.5280 at epoch 1 and then it decreases significantly reaching its lowest point of 0.0050 at epoch 10. This trend is maintained, as the validation loss (orange) is still declining, achieving the value of 0.0010 or less by the epoch 20 and settling at the value of about 0.0005 by the epoch 30. On the other hand, validation loss starts at 0.3920, decreases to 0.2870 at epoch 3 and continues to increase in value to reach 0.5400 at epoch 30. These results indicate that in as much as the model continually decreases the error in the training data, its generalization ability in the validation data decays with time.

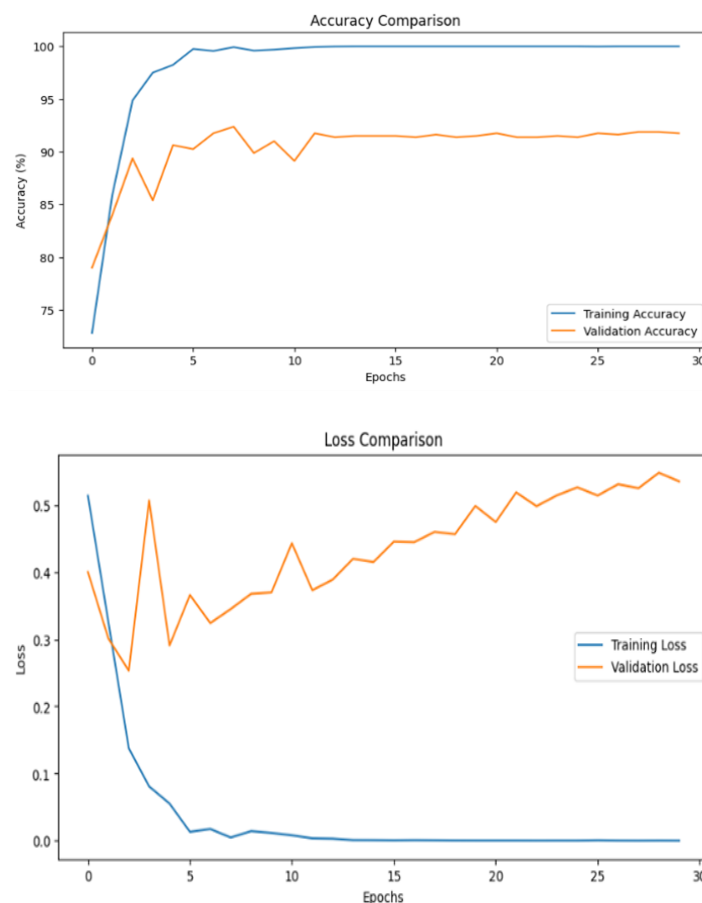


Fig 12. Accuracy& Loss Plot of ConvNeXt.

The present results indicate that even though the training path is stable, expressed in the form of low training loss (approximately equal to 0.0005) and high training accuracy (approximately equal to 0.9980), the accuracy of the considered neural network on the validation set is inclined to stagnate after the first epochs, which leads to a gradual increase in validation loss, which reaches the maximum value of 0.5400. However, the accuracy of validation is still very high, and it varies between 0.9050 and 0.9180 after the 6th epoch, thus proving that the model has a high ability of generalization to the unseen data. All these findings show that despite the model converging quickly and providing decent results on the training set, the training time or adaptive validation monitoring may help further encourage prediction consistency on data in diverse subsets.

Comparison of Models

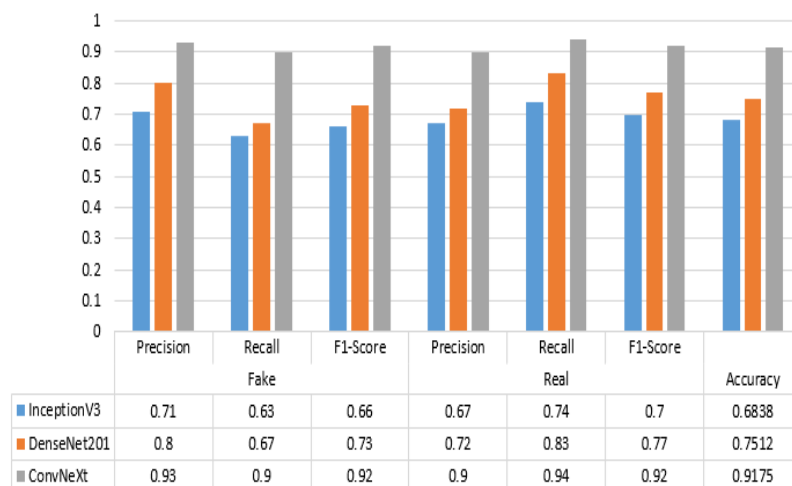


Fig 13. Comparison of Models.

This work examined the performance of various classification models when it is challenged with a corpus that consisted of 400 images that were labelled as Fake and 400 that were labelled as Real. Performance of the models was assessed through precision, recall, F1-score and overall accuracy and each measure was reported individually per class. ConvNeXt was better in all criteria compared to the other approaches. The model achieved a precision of 0.93, recall of 0.90 and F1-score of 0.92 on the Fake class; on the Real class, the precision was 0.89, recall was more than 0.94 and F1-score was 0.92. Therefore, ConvNeXt recorded the best overall accuracy of 91.75 % and consequently provided good generalization between the classes Fig 13 shows Comparison of Models.

DenseNet201 performed on the moderate level. Its specificity to Fake samples was quite high (0.80), but the recall was low (0.67), which means that it was prone to misclassification. The actual recalls were stronger (0.83) and returned an F1-score of 0.77. The model provided an overall accuracy of 75.12 %, which indicated that it was not very effective and could be improved on the aspect of class balance.

The least successful model, InceptionV3, achieved the accuracy of 68.38 %. Detection of Fake content performance was relatively poor, precision of 0.71, recall of 0.63, and F1-score of 0.66, with the Real class metrics being just slightly higher.

Taken together, these findings show that ConvNeXt has a higher discriminatory power that provides high sensitivity and precision in both classes and can be deemed as the most reliable model in detecting fake news in the context of this comparison study.

V. CONCLUSION AND FUTURE SCOPE

This work reveals that detecting manipulated facial images with high accuracy is possible using DenseNet201, InceptionV3 and ConvNeXt transfer learning models. ConvNeXt obtained the highest accuracy of 91.75, plus a precision of 0.93 and recall of 0.90 on fake images and a precision of 0.90 and recall of 0.94 on real images. For DenseNet201, the accuracy recorded was less compared to the ConvNeXt model, indicating small dip in performance precision 0.80 and recall 0.67 for fake and precision 0.72 and recall 0.83 for real. The Accuracy for InceptionV3 is 68.38, with precision 0.71 for fake images and recall 0.63 and precision for real images is 0.71 with recall 0.63. While every model works well for face forgery detection, ConvNeXt gives the best results on both in-domain and out-of-domain data. Moving ahead, further research can add early stopping and look into ensemble learning to improve results for ConvNeXt. Using either small or larger variants of ConvNeXt can make it faster for applications that demand speed. A detection system will benefit from facing adversarial attacks, analyzing worn out or low-resolution inputs and identifying targets using both video and audio. Validating with larger and newer groups of facial images will help construct reliable and adaptable systems against face forgery.

CRediT Author Statement

The authors confirm contribution to the paper as follows:

Conceptualization: Akshatha G, Kempanna M, Ashoka S B and Job Prasanth Kumar Chinta Kunta; **Methodology:** Akshatha G and Kempanna M; **Software:** Ashoka S B and Job Prasanth Kumar Chinta Kunta; **Data Curation:** Akshatha G and Kempanna M; **Writing-Original Draft Preparation:** Akshatha G, Kempanna M, Ashoka S B and Job Prasanth Kumar Chinta Kunta; **Visualization:** Akshatha G and Kempanna M; **Investigation:** Ashoka S B and Job Prasanth Kumar Chinta Kunta; **Supervision:** Akshatha G and Kempanna M; **Validation:** Ashoka S B and Job Prasanth Kumar Chinta Kunta; **Writing- Reviewing and Editing:** Akshatha G, Kempanna M, Ashoka S B and Job Prasanth Kumar Chinta Kunta; All authors reviewed the results and approved the final version of the manuscript.

Declaration of Competing Interest

The authors declare that they have no known competing financial interests or personal relationships that could have appeared to influence the work reported in this paper

Data Availability Statement

The Datasets used and /or analysed during the current study available from the corresponding author on reasonable request.

Funding

No funding agency is associated with this research.

Consent to Publish

All authors gave permission to consent to publish.

References

- [1]. Malik, M. Kuribayashi, S. M. Abdullahi, and A. N. Khan, “DeepFake Detection for Human Face Images and Videos: A Survey,” *IEEE Access*, vol. 10, pp. 18757–18775, 2022, doi: 10.1109/access.2022.3151186.
- [2]. Y. Lu and T. Ebrahimi, “Assessment framework for deepfake detection in real-world situations,” *EURASIP Journal on Image and Video Processing*, vol. 2024, no. 1, Feb. 2024, doi: 10.1186/s13640-024-00621-8.
- [3]. F. Marra, D. Gragnaniello, D. Cozzolino, and L. Verdoliva, “Detection of GAN-Generated Fake Images over Social Networks,” 2018 IEEE Conference on Multimedia Information Processing and Retrieval (MIPR), Apr. 2018, doi: 10.1109/mipr.2018.00084.
- [4]. Z. Liu, X. Qi, and P. H. S. Torr, “Global Texture Enhancement for Fake Face Detection in the Wild,” 2020 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR), pp. 8057–8066, Jun. 2020, doi: 10.1109/cvpr42600.2020.00808.
- [5]. Wang and W. Deng, “Representative Forgery Mining for Fake Face Detection,” 2021 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR), pp. 14918–14927, Jun. 2021, doi: 10.1109/cvpr46437.2021.01468.
- [6]. J. C. Neves, R. Tolosana, R. Vera-Rodriguez, V. Lopes, H. Proenca, and J. Fierrez, “GANprintR: Improved Fakes and Evaluation of the State of the Art in Face Manipulation Detection,” *IEEE Journal of Selected Topics in Signal Processing*, vol. 14, no. 5, pp. 1038–1048, Aug. 2020, doi: 10.1109/jstsp.2020.3007250.
- [7]. J. Yang, S. Xiao, A. Li, G. Lan, and H. Wang, “Detecting fake images by identifying potential texture difference,” *Future Generation Computer Systems*, vol. 125, pp. 127–135, Dec. 2021, doi: 10.1016/j.future.2021.06.043.
- [8]. F. Benchallal, A. Hafiane, N. Ragot, and R. Canals, “ConvNeXt based semi-supervised approach with consistency regularization for weeds classification,” *Expert Systems with Applications*, vol. 239, p. 122222, Apr. 2024, doi: 10.1016/j.eswa.2023.122222.
- [9]. S. C. Hoo, H. Ibrahim, and S. A. Suandi, “ConvFaceNeXt: Lightweight Networks for Face Recognition,” *Mathematics*, vol. 10, no. 19, p. 3592, Oct. 2022, doi: 10.3390/math10193592.
- [10]. Z. Li, T. Gu, B. Li, W. Xu, X. He, and X. Hui, “ConvNeXt-Based Fine-Grained Image Classification and Bilinear Attention Mechanism Model,” *Applied Sciences*, vol. 12, no. 18, p. 9016, Sep. 2022, doi: 10.3390/app12189016.
- [11]. Q. Zhang, X. Wang, M. Zhang, L. Lu, and P. Lv, “Face-Inception-Net for Recognition,” *Electronics*, vol. 13, no. 5, p. 958, Mar. 2024, doi: 10.3390/electronics13050958.
- [12]. H. Qi, R. Han, Y. Shi, and X. Qi, “A Novel High-Performance Face Anti-Spoofing Detection Method,” *IEEE Access*, vol. 12, pp. 67379–67391, 2024, doi: 10.1109/access.2024.3400285.
- [13]. K. Hasan et al., “Facial Expression Based Imagination Index and a Transfer Learning Approach to Detect Deception,” 2019 8th International Conference on Affective Computing and Intelligent Interaction (ACII), pp. 634–640, Sep. 2019, doi: 10.1109/acii.2019.8925473.
- [14]. L. Guamera et al., “The Face Deepfake Detection Challenge,” *Journal of Imaging*, vol. 8, no. 10, p. 263, Sep. 2022, doi: 10.3390/jimaging8100263.
- [15]. M. Elasri, O. Elharrouss, S. Al-Maadeed, and H. Tairi, “Image Generation: A Review,” *Neural Processing Letters*, vol. 54, no. 5, pp. 4609–4646, Mar. 2022, doi: 10.1007/s11063-022-10777-x.
- [16]. Y. Wang, V. Zarghami, and S. Cui, “Fake Face Detection using Local Binary Pattern and Ensemble Modeling,” 2021 IEEE International Conference on Image Processing (ICIP), pp. 3917–3921, Sep. 2021, doi: 10.1109/icip42928.2021.9506460.
- [17]. H. Mittal, M. Saraswat, J. C. Bansal, and A. Nagar, “Fake-Face Image Classification using Improved Quantum-Inspired Evolutionary-based Feature Selection Method,” 2020 IEEE Symposium Series on Computational Intelligence (SSCI), Dec. 2020, doi: 10.1109/ssci47803.2020.9308337.
- [18]. E. Kim and S. Cho, “Exposing Fake Faces Through Deep Neural Networks Combining Content and Trace Feature Extractors,” *IEEE Access*, vol. 9, pp. 123493–123503, 2021, doi: 10.1109/access.2021.3110859.
- [19]. M. Tanaka and H. Kiya, “Fake-image detection with Robust Hashing,” 2021 IEEE 3rd Global Conference on Life Sciences and Technologies (LifeTech), Mar. 2021, doi: 10.1109/lifetech52111.2021.9391842.
- [20]. P. He, H. Li, and H. Wang, “Detection of Fake Images Via The Ensemble of Deep Representations from Multi Color Spaces,” 2019 IEEE International Conference on Image Processing (ICIP), pp. 2299–2303, Sep. 2019, doi: 10.1109/icip.2019.8803740.

Publisher’s note: The publisher remains neutral with regard to jurisdictional claims in published maps and institutional affiliations. The content is solely the responsibility of the authors and does not necessarily reflect the views of the publisher.