



Exploring Xception's Effectiveness in Detecting Deep Fake For Image and Video Formats

Adimalla Rama Rao

Department of Computer Science & Technology, Lendi Institute of Engineering & Technology,
JNTU GV, Vizianagaram, Andhra Pradesh, India ;

* Corresponding Author : Adimalla Ramarao ; ramaraocse501@gmail.com

Abstract: In this paper, a general overview of the use of deep learning in the battle against the booming deepfake industry problem is discussed. Uses of deep learning include, but are not limited to, natural language processing, machine learning, and computer vision and are responsible for a multitude of novel applications. Meanwhile, the growing number of convincing videos and images has also been worrying. This technology can have serious implications online when used for nefarious purposes, such as fake news stories, persona impersonation, financial scams, and distribution of unwanted explicit visual material. Those such as celebrities and politicians are particularly vulnerable to this. The current work evaluates the performance of four deep learning models, namely InceptionResNetV2, VGG19, a standard CNN and Xception, in the fields of deepfakes generation and identification. The evaluation carried out using a Kaggle dataset of deepfakes confirms that Xception outperforms the other models analysed in detecting deepfakes. In a rapidly evolving landscape, where malicious deepfakes are becoming more common every day, it is necessary to develop dependable detection systems to reduce the impact they can have on society.

Keywords: Deep Learning, Deepfake Detection, InceptionResNetV2, VGG19, CNN, Xception.

1. Introduction

With the help of cutting-edge machine learning capabilities, deepfake tech has advanced to the point that it can generate very convincing fake videos and pictures by superimposing a face or image from one person onto another. This has raised concerns about the potential misuse of such content, notably to disseminate misinformation or to influence public opinion in negative ways. However, with an increasing threat of outbreaks, scientists have turned to deep learning to create detection tools [1]. A subset of state-of-the-art neural networks, convolutional and recurrent neural networks have been shown to detect the negligible sleight made by deepfakes. Through pattern recognition and by observing a vast amount of data, these networks can identify fraudulent image compared with the original. Several previous works have contributed to the progress of deepfake detection methods. For example, H. Li et al. used movement of facial muscles to identify fraudulent expressions in video (Li et al., 2020) and J. Rossler et al. looked at micro features like head motion and blink patterns to determine indications of

manipulation (Rossler et al., 2019). The Introduction outlines the need for a problem and highlights the need for deepfake detection research to be topical. The next sections detail specific techniques, advances and challenges in identifying manipulated media and explain why there is an effort to maintain trust [2] in digital content in the face of increased sophistication in manipulation by AI.

The focus of this work is the solution of the growing challenge of deepfake videos and images by leveraging deep learning. The study compares InceptionResNetV2, VGG19, CNN and Xception models on a Kaggle dataset of deep-fake images, ensuring that reliable deep fake detection tools are essential in an era where deep-fakes can be increasingly leveraged for fake news [3], defrauding systems and violating privacy. Deepfakes are synthetic but believable videos and images that raise a significant threat in a number of fields, including false information and privacy. In this paper, the impact and resistance to these 'bad uses' will be explored, highlighting the need for a more robust detection strategy.



2. Literature Survey

However, with the current popularity of deepfake technology, people have also been worried that this data can be misused, which has led to detection being a research category in gathering. With the current popularization of deepfake, however, people have also been concerned that such data can be misused, making detecting deepfakes an increasing research category [4]. Much of the advancement in this area is founded on established datasets like DeepfakeDetection and FaceForensics++, but the vast majority of these are videos captured by volunteers in studio environments and thus lack the broadness to capture the full spectrum of deepfakes in their true forms that are generally found online. To seek a solution to this, the authors suggest the WildDeepfake dataset, consisting of 7,314 face sequences sampled from 707 deepfake videos from the web.

WildDeepfake is designed to capture the diversity and chaos of authentic deepfake material, in contrast to previous benchmarks. The composition makes it more difficult for detection algorithms, as it differs from the "normal" conditions and generation tools used for older data sets [5]. The authors compare the performance of various baseline detection networks to conventional data sets and WildDeepfake, and find that their performance decreases significantly with WildDeepfake. To deal with this, they suggest two Attention-based Deepfake Detection Networks (ADDNets), which utilize an interest masks over valid and faked faces [6]. These networks did not only achieve good performance on traditional benchmarks but also on the more challenging set WildDeepfake, making them potentially valuable in combating actual deepfake attacks. In total, this makes the work an important contribution towards developing detection systems capable of keeping up with such deepfake material in the wild. This work investigates a novel light on deepfake detection using the consistency of the source features in the manipulated images. What this really comes down to is that there is still a faint hint of the original source even after they've gone through complicated deep fake creation [7]. For this they develop pair-wise self-consistency learning (PCL), a representation-learning approach through Convolutional Neuronal Networks which can isolate these unknown source features as indications of manipulation.

Together with PCL, they release the inconsistency image generator (I2G), a synthesis technique that generates many samples heavily annotated for PCL. The models generated from the datasets are tested nine different ones to demonstrate their improvement in detection performance; those are widely used datasets [8]. Individually the average AUC increases from 96.45% to 98.05% which

surpasses former best state of the art results, and in cross-dataset testing from 86.03% to 92.18%. The results demonstrate the effectiveness of the proposed method in enhancing detection accuracy and reliability for various types of data sets, providing valuable methodology and understanding for further research on the efforts to protect against deepfake content [9]. This paper proposes a novel detection method based on Attribution-Based Confidence (ABC) metric to overcome the increasing danger of AI-generated fake videos from powerful Generative Adversarial Networks (GANs). Notably, this method does not require access to training data and there is no calibration step on the validation data: inference can be done only given the trained model. In practice the authors simply only train a model with real videos and use the ABC metric to determine whether a new video is real or not. The metric gives a score and real videos are found to score above a threshold of 0.94 on this metric [10]. The ABC metric is suitable for efficient discrimination between real videos and manipulated ones as it is not vulnerable to extensive training data and validation data.

Overall, the work demonstrates that using the attribution approach, one can effectively distinguish between original and manipulated (deepfake) video content, providing another possibly applicable method in the fight against the increasingly sophisticated video manipulation technologies [11]. This paper highlights an important vulnerability of existing deepfake detectors: adversarial perturbations can trick deepfake detectors. The authors create perturbed deepfake images in both blackbox and whitebox settings via Fast Gradient Sign Method and Carlini-Wagner L2 attack, respectively, which drastically degrade detection performance. However, it was revealed that accuracy of the detectors below 27% for perturbed images, while they have outperformed ordinary deepfakes above 95% accuracy. Two strategies are then described for enhancing detector robustness followed by the examination of the paper.

The first technique involved in Lipschitz regularization restricts the degree of change in the detector's output when the input undergoes slight alterations, thereby improving the blackbox accuracy of deepfakes perturbed by noise by approximately 10%. The second, called the Deep Image Prior (DIP) defense, removes perturbations on unsupervised examples by using generative convolutional networks that successfully removed perturbations from perturbed deepfakes that the detector incorrectly classified, yet maintained an accuracy of 98% on a 100-image subset of unperturbed cases [12]. The paper points out that current deepfake detectors are easily tricked by adversarial manipulation methods and provides practical approaches to making these methods more robust. With the advent of Generative Adversarial Networks and their ability to generate more and more convincing

video content, the paper proposes a deep learning-based spotting framework that combines multiple deepfakes to boost the accuracy of video manipulation detection. As the authors point out, deepfakes are already being employed for propaganda, cybercrime and meddling in political conversation, making effective countermeasures critical. Inspired by the recent success of deep learning in the field of classification [13], DeepfakeStack employs multiple “fresh” deep learning-based classification models to enhance classifiers' performance in the detection of deep learning-based fake images. In testing, this ensemble achieved the highest accuracy of 99.65% and AUROC value of 1.0 when compared with individual classifiers. The findings indicate that DeepfakeStack can be a valuable starting point for developing robust deepfake detection algorithms for real-time use, offering a comprehensive defense against the misuse of hyper-realistic synthetic media [14]. The paper makes a valuable contribution to an active interest in technologies that would be able to keep up with the rapidly changing landscape of manipulated audio and video content.

3. Related Work and Methodology

3.1. Proposed Work

All this is what the system suggested here aims to do to tackle the problematic deepfakes of these days by incorporating multiple deep learning models in one evaluation pipeline. It is built on top of Kaggle deepfake dataset and is testing the performance of the detection model InceptionResNetV2, VGG19 and a standard CNN simultaneously [14]. The central aim is to be able to accurately differentiate between authentic and manipulated content. The details of the deepfakes' construction are being explored to help create a better understanding regarding the way that they are created and disseminated. Multiple algorithms run concurrently can be compared more completely and useful training and testing the variety of the dataset is appropriate. As these two teams learn, become more accurate and reliable in their ability to accurately detect a deepfake helps curb the effects of fake news, impersonation [15], and privacy infringement. The proposed system is a practical solution in this regard to counter the detrimental utilization of Deepfakes in the cyber space.

3.2. System Architecture

Deepfakes are detected as follows: a multi-stage pipeline involving the following steps is followed: imports of deepfake videos, exploratory data analysis (EDA) and data visualization, resizing pictures, and the use of an image data generator. After the EDA stage, the system applies a few deep learning models such as InceptionResNetV2, VGG19, CNN, and Xception to perform the detection process [16]. The first step is that the deepfake videos are all uploaded and then broken down by frame, providing a more granular look at the individual frames and the

potential to manipulate them at that level. In EDA, visualization is applied to get patterns in the frame data to the surface and to create a more complete picture of the whole data set. Images are then resized to account for different frame sizes, which will cause the model to “generalize”, and an additional image generator will add variation to make the model more robust [17]. The core of the pipeline are the four deep-learning models InceptionResNetV2, VGG19, CNN and Xception, which are well known for image-classification capabilities [18]. Each model works independently of each other to extract unique features from the frames and contributes to the final decision of detection. Last but not least, the models are evaluated in terms of accuracy, precision, recall, F1-score, specificity, sensitivity, MAE and MSE. In conclusion, this pipeline provides a systematic and effective approach to tackling deepfake detection, incorporating both preprocessing and exploratory analysis stages before entering into the realm of cutting-edge deep learning techniques.

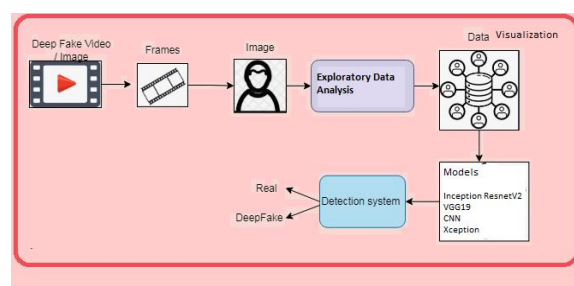


Fig. 1 System Architecture

3.3. Dataset Collection

Creating a robust deepfake detection system is crucial and requires a comprehensive dataset that contains a wide range of authentic and fake video clips [18]. A deepfake dataset for Kaggle is just that it offers a massive and diverse collection of real and manipulated videos. Kaggle, a leading platform for data science and machine learning competitions, offers such collections readily, which can be applied to solve problems related to deep fake. The data is started from the collection of deepfakes from kaggle that typically contain shed-tons of verified and fake footage. Careful labelling of each video, so that the data set is appropriately labelled for supervised learning [19]. The range of actors, settings, and contexts of the most diverse type of scenes that the collection covers will help to ensure that the resulting model is not overly sensitive to any specific type of scene. As it is a collaborative platform, rates of sharing on Kaggle also typically include methods to tidy and improve datasets; the resulting overall improvement of quality and utility of data is useful for research purposes. Accessing reliable labels can also allow researchers to properly train and evaluate detection models, with their results being compared to a similar standard.

3.4. Image Processing

As for computer vision, correctly preparing the images is an important step to emphasize when building strong computer vision models, and this is also the case with deepfake detection. Here two techniques are being used; image resizing and an image data generator. Resizing reduces the image sizes to a standard size, ensuring that all images fed into the model are in a similar format. It is important because it avoids redundant computations for “noise” in the resolution and enables seamless integration of images into a neural-network architecture without losing too much of the picture fine details while maintaining sufficient efficiency [20]. The image data generator, on the other hand, attempts to add some variation to the dataset through random transformations like rotations, zooms and horizontal flips which makes the model robust enough to deal with any variation in reality without relying on few examples. For the deepfake detection task, the augmentation step is significant because, in reality, video frames can be augmented in numerous ways. Resizing normalizes inputs and maintains consistency across them, while the additional variation by the data generator introduces useful variation and helps provide a model that can be used to distinguish between actual and manipulated inputs. These methods are complementary; the visual information they will face when put into practice in real life will be of varying types and complexity.

3.5. Data Visualization

A crucial piece of the puzzle in the fight against deepfakes is analyzing videos frame-by-frame, as changes can be observed and examined over time. After splitting video into frames, visualization is used to look for patterns in the new video frame set. This can come in many shapes histograms of pixel intensity, distributions of frame durations, or optical flow mapping of the content over time all of which provide insight into behaviour of data over time [21]. For example, the data in a histogram might reveal rare, unexpected changes in pixel intensity that could be signs that the image has been manipulated. Side-by-side comparison of frames, or creating heat maps of motion, can also reveal irregularity that occurs in the deepfake creation process. To introduce further insight into optical flow, you'll draw vectors that represent the movement of the scene between two consecutive frames, called optical flow vectors. The ability to “zoom in” to any individual frame and mark anything suspicious on it through interactive visualization tools also makes it easy to find out more detailed information. These approaches can work together to identify patterns, anomalies and contaminations within the temporal data that the researcher can use to improve the description of the detection models. Finally, this type of visualization will help in decision making during the development of the

model and will assist in providing an easier to understand detection system.

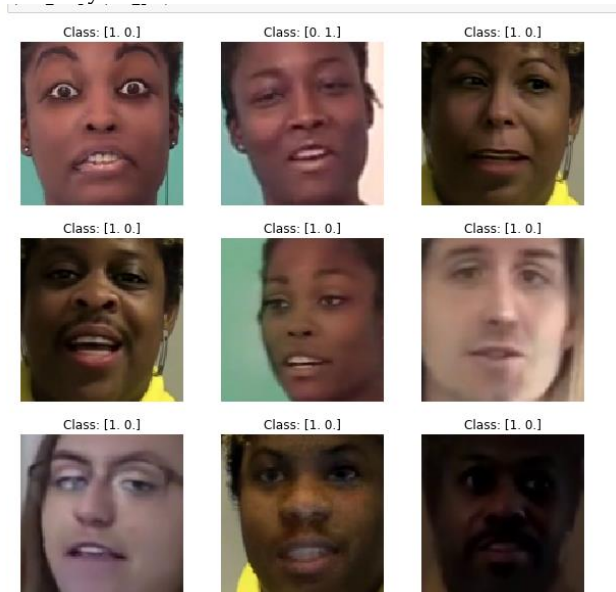


Fig. 2 Data Visualization

3.6. Algorithms

InceptionResNetV2: This is the combination of two architectures namely Inception and ResNet, which makes it a deep network fitted to extract rich features. It is good at identifying subtle and layered patterns, making it effective at detecting visual clues that can be left in the wake of a deepfake.

VGG19: The model VGG19 was chosen because of its simplicity and trustworthiness. It has an architectural structure that includes small receptive field across the board that allows it to capture both low-level and high-level visual features. This works well to detect patterns that can be associated with any particular frame being a falsified copy of some other frame caused by deepfake manipulation.

CNN: A CNN is added as a means of extracting spatial features from image data, which is known for its ability to extract spatial features of images. It is relatively light weight, making it an ideal choice for speedy frame-by-frame processing for detection.

Xception: Xception is selected due to its extensive features and robust feature-extraction performance, especially its capability for modelling the spatial hierarchy in an image. This is a deep model for picking up complicated patterns related to manipulation suitable for this project.

4. Results and Discussions

Accuracy: The accuracy of a test indicates how well the test discriminates between the two outcome groups in the correct direction. Defined as the percentage of correct responses (true answers) in the entire population of respondents taken; expressed as:

$$\text{Accuracy} = \frac{TP + TN}{TP + TN + FP + FN}.$$

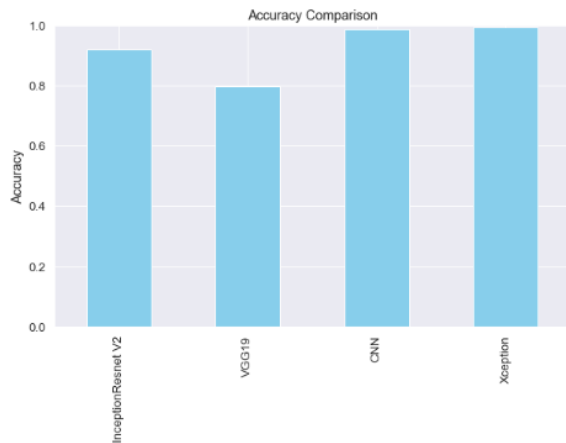


Fig. 3 Accuracy comparison graph

Precision: Percentage of true positive values. Applying the formula, it is computed as:

$$\text{Precision} = \frac{\text{True positives}}{\text{True positives} + \text{False positives}} = \frac{TP}{TP + FP}$$

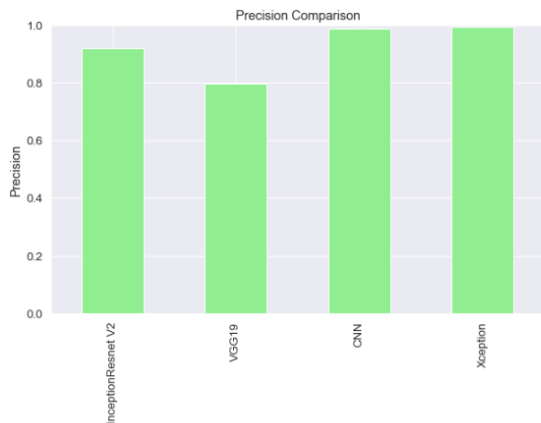


Fig. 4 Precision comparison graph

Recall: In machine learning, recall is the metric that is used to determine how well a machine learning model can be used to identify all items associated with a specific class. It is the ratio correctly identified positive / actual positives it gives an idea of what is the amount of true expressed by the model within a certain class.

$$\text{Recall} = \frac{TP}{TP + FN}$$

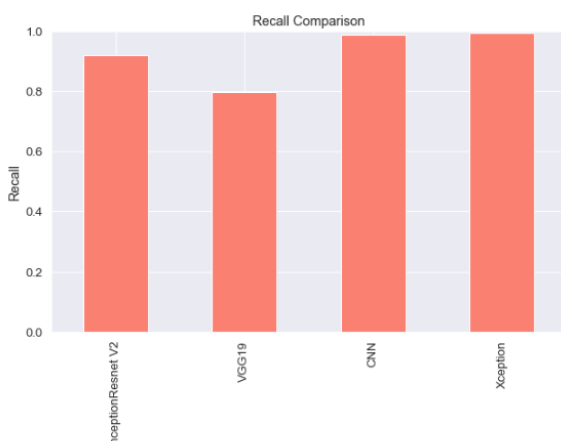


Fig. 5 Recall comparison graph

F1-Score: F1-score is a metric which combines precision and recall in one single number describing overall model accuracy. Measures the accuracy of the predictions made by the model over the entire dat

$$\text{F1 Score} = \frac{2}{\left(\frac{1}{\text{Precision}} + \frac{1}{\text{Recall}}\right)}$$

$$\text{F1 Score} = \frac{2 \times \text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}}$$

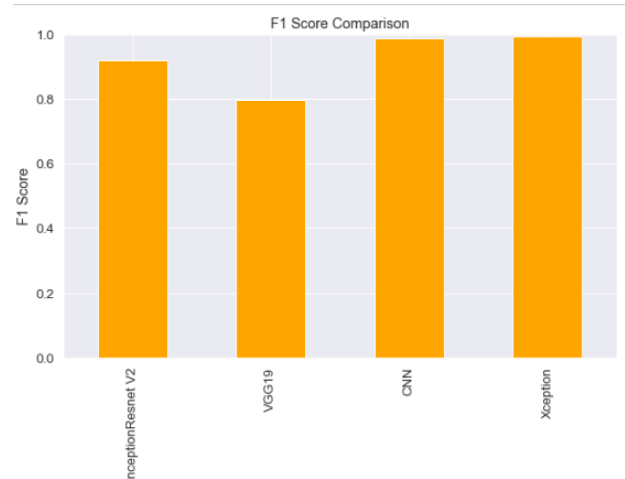


Fig. 6 F1 Score comparison graph

Sensitivity: The sensitivity of a test indicates how well it detects a positive test result. It is the percentage of true positives (TP) out of all positive cases,

$$\text{Sensitivity} = \frac{TP}{TP + FN}$$

$$\text{Sensitivity} = \frac{\text{True Positives}}{\text{True Positives} + \text{False Negatives}}$$

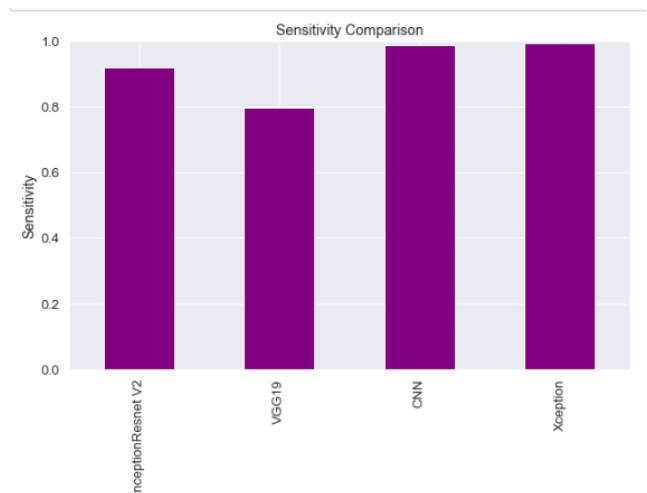


Fig. 7 Sensitivity comparison graph

Specificity: Specificity is its accuracy in identifying negative cases. The calculation of it is the ratio of number of true negatives to the number of actual negatives cases, and is expressed as:

$$\text{Specificity} = \frac{TN}{TN + FP}$$

$$\text{Specificity} = \frac{\text{True Negatives}}{\text{True Negatives} + \text{False Positives}}$$

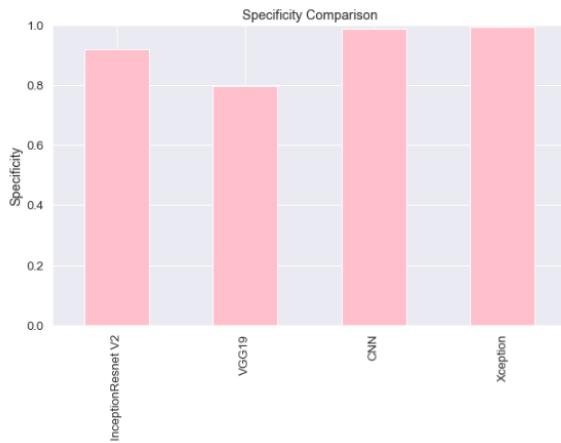


Fig. 8 Specificity comparison graph

MAE: MAE is the sum of the absolute error of all the samples divided by n , the number of samples, which would give a mean value without regard for its direction; this value indicates the average absolute error regardless of direction from the predicted and true values. Relative frequencies are used to weight some of the formulations.

$$MAE = \frac{1}{n} \sum \left| y - \hat{y} \right|$$

Divide by the total number of data points

Predicted output value

Actual output value

Sum of

The absolute value of the residual

MSE: Mean squared error (MSE) is the average of the squared differences between observed values and the values that a model predicts. The point of squaring the differences is that the raw differences can be positive or negative. If a prediction is greater than the actual value, then the difference is positive, and if it is less than the actual value, the difference is negative. When the raw differences are not squared, positive and negative differences will tend to cancel each other out. Squaring corrects this problem, and, when squared, the error always indicates the size of the difference. The formula for the mean squared error is

$$MSE = \frac{1}{n} \sum_{i=1}^n (Y_i - \hat{Y}_i)^2$$

MSE = mean squared error

n = number of data points

Y_i = observed values

$\hat{Y}_i$ = predicted values

	Accuracy	Recall	Precision	F1 Score	Sensitivity
InceptionResNet V2	0.918432	0.918415	0.918415	0.918415	0.918415
VGG19	0.796416	0.796435	0.796435	0.796435	0.796435
CNN	0.987048	0.987036	0.987036	0.987036	0.987036
Xception	0.993900	0.993902	0.993902	0.993902	0.993902

	Specificity	MAE
InceptionResNet V2	0.918415	<function mae at 0x0000204027E2288>
VGG19	0.796435	0.32427
CNN	0.987036	0.019514
Xception	0.993902	0.009542

	MSE
InceptionResNet V2	<function mse at 0x0000204027E23A8>
VGG19	0.162385
CNN	0.009877
Xception	0.004718

Fig. 9 Performance evaluation table

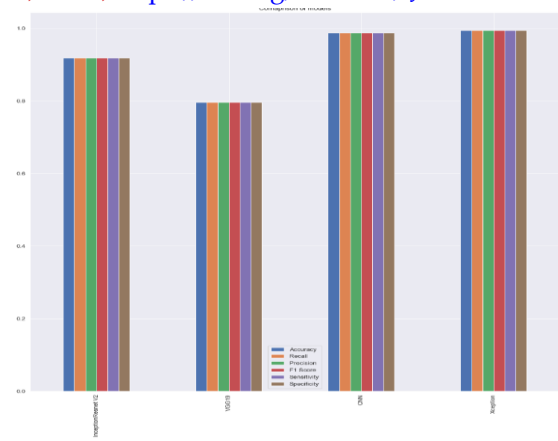


Fig. 10 Performance evaluation comparison graph

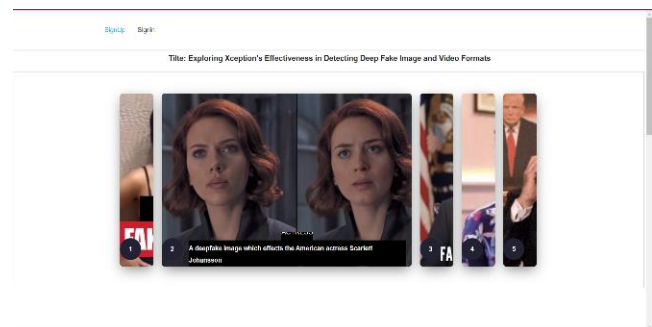


Fig. 11 Home & about page

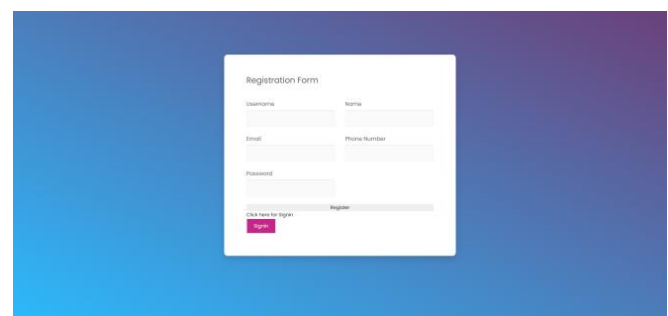


Fig.12 Registration page

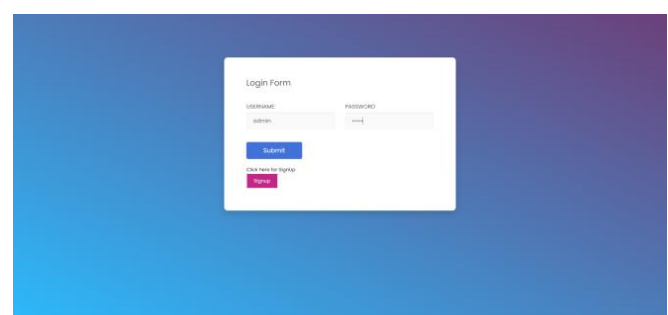


Fig. 13 Login page

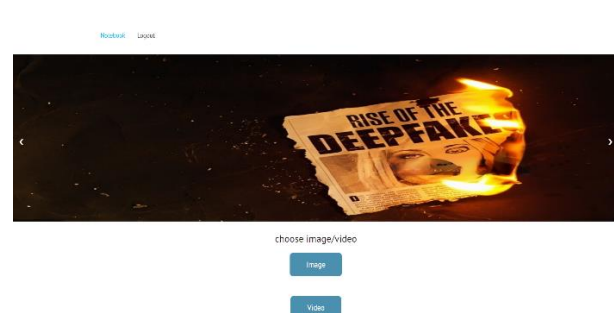


Fig. 14 Main page



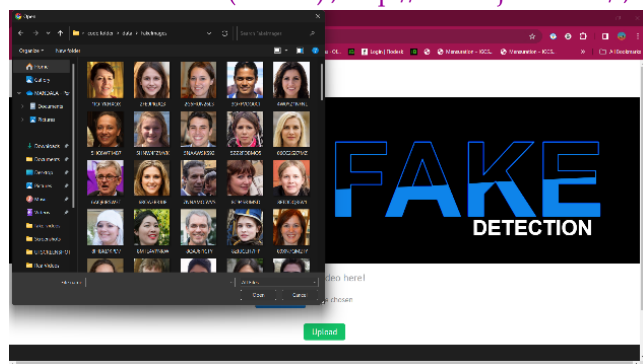


Fig. 15 Upload input

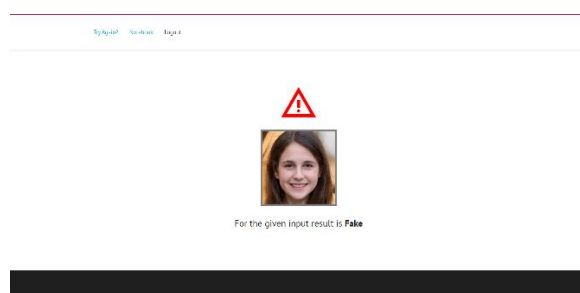


Fig.16 Predict result as fake

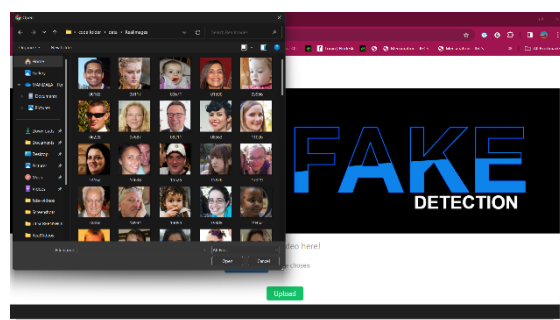


Fig. 17 Upload another input

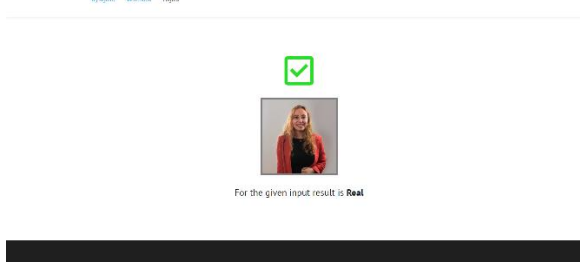


Fig.18 Final outcome as real for given input

5. Conclusion and Future Work

Overall, this study just serves to further emphasize the need to continuously enhance deepfake detection systems as the manipulation techniques continue to become increasingly sophisticated. After examining older and more recent strategies, one can see that there are actual shortcomings in adjusting to new threats and identifying mangled patterns. Therefore the presented system is a more proactive one and is based on the novel models such as InceptionResNetV2 and Xception. By subjecting the

system to a vast array of scenarios on a deepfake data set, training will provide exposure in order to prepare the system for real-world conditions. In addition to enhancing raw detection capabilities, the strategy seeks to enhance understanding of deepfakes and their potential to cause widespread harm to society—such as misinformation, privacy breaches, or impersonation attacks—as well as expand the amount of legal and educational avenues for tackling their misuse. Beyond raw detection enhancements, the approach also aims to deepen the understanding of deepfakes and their potential to inflict harm upon society, including misinformation, privacy leaks, and impersonation attacks, while also introducing new legal and educational opportunities for combating the misuse of deepfakes. That's certainly a positive development for a safer and more secure online world for all.

In the future, this area of research can be extended based on emerging new threats from deepfakes. Other algorithms or newer techniques (e.g. reinforcement learning) may result in further improvement in detection performance. Cooperation among academia, industry and policymakers will be likely key to progress in the development of common standards for deepfakes detection. It may also be of interest to future work to develop real-time detection systems and introduce explainability features which allow to better grasp and trust the detection decision.

References

- [1]. M. Mirza and S. Osindero, "Conditional generative adversarial nets," arXiv preprint arXiv:1411.1784, 2014.
- [2]. Y. Bengio, P. Simard, and P. Frasconi, "Long short-term memory," IEEE Trans. Neural Netw., vol. 5, pp. 157–166, 1994.
- [3]. I. Goodfellow, Y. Bengio, and A. Courville, Deep learning. MIT press, 2016.
- [4]. S. Hochreiter, "Ja1 4 rgenschmidhuber (1997). "long short-term memory", " Neural Computation, vol. 9, no. 8.
- [5]. M. Schuster and K. Paliwal, "Networks bidirectional recurrent neural," IEEE Trans Signal Proces, vol. 45, pp. 2673–2681, 1997.
- [6]. J. Hopfield et al., "Rigorous bounds on the storage capacity of the dilute hopfield model," Proceedings of the National Academy of Sciences, vol. 79, pp. 2554–2558, 1982.
- [7]. Y. Wu, M. Schuster, Z. Chen, Q. V. Le, M. Norouzi, W. Macherey, M. Krikun, Y. Cao, Q. Gao, K. Macherey, et al., "Google's neural machine translation system: Bridging the gap between human and machine translation," arXiv preprint arXiv:1609.08144, 2016.
- [8]. R. K. S, T. Y, and S. C. U, "Implementation of Pulmonary Disease Detection Model using Deep Learning," International Journal of Research and Development in Engineering Sciences, vol. 7, no. 2, p. 19, Mar. 2025, doi: 10.63328/ijrdes-v7ri2p4.
- [9]. Ravi Samikannu, Retrieval-Augmented Multi-Agent Architecture for Hallucination Reduction in Large Language Models, International Journal of Computer Science , Engineering and Artificial Intelligence , Vol. 2, Issue. 4 (October – December) , 2025 , Pages: 7 – 13. DOI : <https://doi.org/10.63328/IJCSEAI-V2RI4P2>
- [10]. Ujval Mutyala , " Embodied Agentic AI for Autonomous Task Planning and Execution in Dynamic Environments ", International Journal of Computer Science , Engineering and Artificial Intelligence , Vol. 2, Issue. 4 (October – December) , 2025 , Pages: 22 – 30. DOI : <https://doi.org/10.63328/IJCSEAI-V2RI4P4>

- [11]. Naseer, I., Akram, S., Masood, T., Rashid, M., & Jaffar, A. (2023). Lung cancer classification using modified U-Net based lobe segmentation and nodule detection.
- [12]. Pan, S. J., & Yang, Q. (2010). A survey on transfer learning. *IEEE Transactions on Knowledge and Data Engineering*, 22(10), 1345–1359.
- [13]. Pehrson, L. M., Nielsen, M. B., & Lauridsen, C. A. (2019). Automatic pulmonary nodule detection applying deep learning or machine learning algorithms to the LIDC-IDRI database: A systematic review.
- [14]. Perez, L., & Wang, J. (2017). The effectiveness of data augmentation in image classification using deep learning, *arXiv*. <https://doi.org/10.48550/arXiv.1712.04621>.
- [15]. Raza, R., Zulfiqar, F., Khan, M. O., Arif, M., Alvi, A., Iftikhar, M. A., & Alam, T. (2023). Lung-EffNet: Lung cancer classification using EfficientNet from CT-scan images.
- [16]. Shin, H. C., Roth, H. R., Gao, M., Lu, L., Xu, Z., Nogues, I., Yao, J., Mollura, D., & Summers, R. M. (2016). Deep convolutional neural networks for computer-aided detection: CNN architectures, dataset characteristics and transfer learning. *IEEE Transactions on Medical Imaging*, 35(5), 1285–1298.
- [17]. Shorten, C., & Khoshgoftaar, T. M. (2019). A survey on image data augmentation for deep learning. *Journal of Big Data*, 6(1).
- [18]. Sokolova, M., & Lapalme, G. (2009). Beyond accuracy, precision, recall and F-measure: A review of evaluation metrics.
- [19]. Yang, X., Duan, A., Jiang, Z., Li, X., Wang, C., Wang, J., & Zhou, J. (2026). Segmentation and classification of lung cancer images using deep learning. *Applied Sciences*, 16(2), 628. <https://doi.org/10.3390/app16020628>.
- [20]. Yogesh Kumaran, S., Jospin, J. J., Mahesh, T. R., Bhatia, S., Alzahrani, S., & Alojail, M. (2024). Explainable lung cancer classification with ensemble transfer learning of VGG16, ResNet50 and InceptionV3 using Grad-CAM.
- [21]. Zhi, L., Jiang, W., Zhang, S., & Zhou, T. (2023). Deep neural network pulmonary nodule segmentation methods for CT images: Literature review and experimental

How To Cite:

Adimalla Rama Rao , “ Exploring Xceptions's Effectiveness in Detecting Deep Fake For Image and Video Formats “, *International Journal of Computational Sciences and Engineering Research (IJCSER)* , Volume 3, Issue 3 , pp 19 - 27 , July 2026. CrossRef DOI: <https://doi.org/10.63328/ijcser-v3ri3p3>

Declaration

Conflicts of Interest: The authors declare no conflict of interest.

Author Contribution: All authors wrote the main manuscript text and also consent to the submission.

Ethical approval: Not applicable.

Consent to Participate: All authors consent to participate.

Funding: Not applicable, and No funding was received

Institutional Review Board Statement: Not applicable.

Informed Consent Statement: Not applicable.

Personal Statement: We declare with our best of knowledge that this research work is purely Original Work and No third party material used in this article drafting. If any such kind material found in further online publication, we are responsible only for any judicial and copyright issues.

Acknowledgements

We thank everyone who inspired our work.

:: Overview of the Paper for Researchers ::

