

Video Forgery Detection Using Machine Learning

Nerella Bharath Ganesh Dept of ECE IARE

Dr. S China Venkateshwarlu Professor Dept of ECE IARE

Dr. V Siva Nagaraju Professor Dept of ECE IARE

Abstract - With the rapid advancement of digital media editing tools, video forgeries have become increasingly sophisticated and harder to detect, posing serious threats to information security, legal evidence integrity, and digital content authenticity. This project presents a machine learning-based approach for the detection of video forgeries, focusing on identifying tampered frames, splicing, and copy-move operations. The system extracts spatiotemporal features and motion inconsistencies using techniques such as optical flow analysis, frame differencing, and noise residual examination. These features are then analyzed using supervised machine learning models, including Convolutional Neural Networks (CNNs) and Support Vector Machines (SVMs), to distinguish.

Key Words: Video Forgery Detection, Machine Learning, Digital Forensics, Frame Tampering, Copy-Move Forgery, Splicing Detection, Convolutional Neural Networks (CNN), Support Vector Machine (SVM), Optical Flow, Feature Extraction, Deep Learning, Video Manipulation, Authenticity Verification.

1.Introduction

With the rapid growth of digital media and video-sharing platforms, ensuring the authenticity of video content has become increasingly challenging. Forged videos, often indistinguishable from genuine ones to the naked eye, can be used maliciously in political manipulation, cybercrime, and legal tampering. Traditional manual methods for detecting such forgeries are time-consuming and prone to error. As a result, automated video forgery detection using machine learning techniques has emerged as a promising solution.

This paper proposes a system that uses machine learning to detect common types of video forgeries such as splicing, copy-move, and frame duplication. By extracting spatial and temporal features from video frames, the system leverages supervised learning models to classify whether a video has been tampered with. everyone can benefit from the clarity and structure that a well-designed task manager provides.

2. Body of Paper

The video forgery detection system is developed using a Python-based machine learning architecture. The system employs a Jupyter Notebook environment for development, Python as the primary programming language, and various machine learning libraries to implement core detection algorithms. The system adheres to a modular design principle, separating data preprocessing, model training, and result visualization to ensure maintainability and adaptability.

- **Jupyter Notebook:** Provides an interactive environment for code development, experimentation, and visualization.
- **Python:** The core programming language used for implementing algorithms and data handling.
- **Machine Learning Libraries (e.g., TensorFlow/PyTorch, OpenCV, scikit-learn):** Provide essential tools for image/video processing, feature extraction, and model development.
- **Detection Algorithms:** Implement the core logic for identifying inconsistencies and anomalies indicative of video forgery.

The application follows a client-server model:

- **Client (Frontend):** Users interact via a simple HTML/CSS-based interface. Optionally, JavaScript can enhance interactivity.
- **Server (Backend):** Java Servlets manage request routing, task processing, and database communication.
- **Database:** MySQL (via XAMPP) stores user data, task lists, statuses, deadlines, and metadata.

Table -1:

Year	Study/Project	Summary
2021	Li et al. – Detecting Video Forgery with Compression Artifacts	Explored the use of compression artifacts and their inconsistencies across frames as indicators of video manipulation.
2022	Haliem et al. – Splicing Detection using Noise Features	Proposed a method for detecting video splicing by analyzing variations in noise patterns (PRNU) in different regions of video frames.

2023	Zhao et al. – Multi-Modal Fusion for Deepfake Detection	Investigated combining visual, audio, and physiological signals (e.g., heart rate from face) for more robust deepfake detection.
------	---	--

- Detection/Inference: Applying the trained model to new videos.
- Result Visualization: Presenting detection outcomes in an interpretable manner. This.

3. SYSTEM ARCHITECTURE

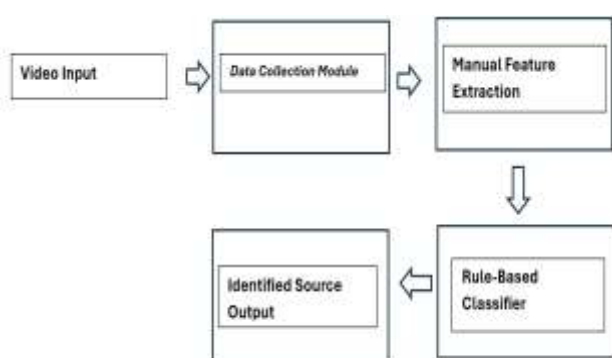
The proposed system is to develop an automatic system that has a built-in classification model and a segmentation model. The first part of building the model is pre-processing and then split into training, testing, and validation data. The data is further trained using multiple models and their respective results are compared after which the best model is used for further classifications. Similar procedure is carried out for building the segmentation model.

The proposed methodology includes following tasks, in essence, dataset splitting, dataset transformation, hyperparameter tuning, modeling. Further details of each task are discussed below.

1) DATA SPLITTING:

In general, there are various ways of splitting the data based on the given data. The dataset is spitted using hold-out technique, in essence 85% of training data and 15% of testing or validation data, considering the fact that the data set used in this study isn't biased and has almost equal amount of classes. Also, we have not used

Existing Block Diagram



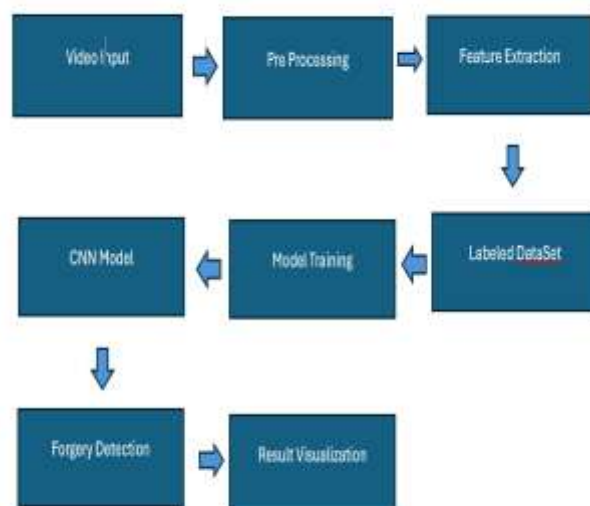
Proposed Block Diagram

Fig -1: Figure

While the detection process is automated, the presentation of results to human analysts is crucial. The system design should follow HCI principles, focusing on clear visualization of detected anomalies, confidence scores, and explanations for the detection, allowing forensic experts to interpret and validate the findings efficiently.

Software Architecture and Modular Design The project adopts a modular architecture, breaking down the complex problem into manageable components:

- Data Acquisition: Handling various video formats and sources.
- Preprocessing: Normalizing video streams, frame extraction.
- Feature Extraction: Deriving forensic features or feeding raw data to deep learning models.
- Model Training: Training machine learning models on labeled datasets.



any stratification algorithms.

2) PREPROCESSING OF DATA:

Digital pictures vary significantly in terms of picture size. Hence all the images are scaled to a common frame of reference depending upon the model we are using. As our emphasis is majorly on the brain region, unnecessary regions of the skull are removed (a.k.a., Skull 14 Stripping). In the end, the dataset is

normalized using the Z-score. Once these steps are accomplished, the dataset is split into training, testing, and validation.

3)BUILDING DIFFERENT CLASSIFICATION MODELS:

Two deep learning models are built using VGG16 [20] and ResNet50. Although, the base model was built using the multilayer perceptron and used transfer learning in order to increase the performance of the model. The models are built by taking care of overfitting and underfitting cases while they are monitored and are evaluated using classification evaluation metrics. Early Stopping and loss monitoring call-backs were also used while training the model.

4) SEGMENTATION MODELING:

In general, the Unet architecture proposed by Olaf Ronneberger is proven the best for segmentation tasks. The residual blocks in the unet architecture helps the CNN model to get significant results. In this study, we proposed a hybrid architecture blending the residual blocks and the Unet architecture and we call it ResUnet. There are various evaluation metrics for the classification model. Here are the few that we have considered. In this study we have used Tversky loss function in order to make evaluation more fault free. This loss function has parameters that can be optimized to get significant results by penalising the loss function. This loss function contains constants 'alpha' and 'beta' to act as penalising factors in case of false positives and false negatives. In our study we have used

$$\alpha = 0.7$$

1)Accuracy:

Accuracy refers to the ratio of the instances that are correctly classified to the total number of instances. It is used to evaluate the classification model when the data is balanced.

Here's the formula $\text{Accuracy} = (\text{TP} + \text{TN}) / (\text{TP} + \text{TN} + \text{FP} + \text{FN})$
 FN-False Negative,
 FP- False Positive,
 TN- True Negative,
 TP- True Positive.

2) Recall:

Recall refers to the ratio of true positives to the sum of true positives and false negatives. This evaluation metric is used as we want to check the efficiency of false negatives. Here's the formula

$$\text{Recall} = \text{TP} / (\text{TP} + \text{FN})$$

3)Confusion Matrix :

A Confusion matrix is a 2 X 2 matrix which quantifies about all the possible outcomes of the classification model, in essence, the true positive, true negative, false positive, and false negative.

Implementation Steps for implementation

1.Install Required Software & Tools Install Required Software & Tools

2 Set Up a Virtual Environment

3 Install Dependencies (tensorflow keras opencv python matplotlib pandas) .

4 Download & Preprocess the Dataset.

5. Train the Dataset.

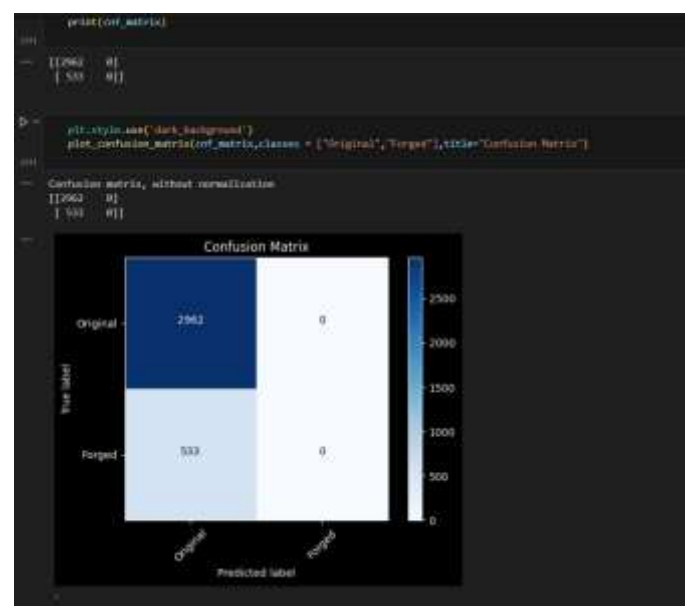
6.Test the Model with Input Vidoes.

7. Run Python file for Execution.

Result



Confusion Matrix



REFERENCES

1. Ahmed, F., Khelifi, F., Lawgaly, A., Bouridane, A., 2019. Comparative analysis of a deep convolutional neural network for source camera identification. In: 2019 IEEE 12th International Conference on Global Security, Safety and Sustainability (ICGS3).
2. IEEE, pp. 1–6.
3. Aizvorski, 2014. h264bitstream. URL: <https://github.com/aizvorski/h264bitstream>.
4. (Accessed 15 July 2024).
5. Akbari, Y., Almaadeed, N., Al-Maadeed, S., Khelifi, F., Bouridane, A., 2022. Prnu-net: a deep learning approach for source camera model identification based on videos taken with smartphone. In: 2022 26th International Conference on Pattern Recognition (ICPR), IEEE, pp. 599–605.
6. Altinisik, E., Sencar, H.T., 2020. Source camera verification for strongly stabilized
7. Dealing with video source identification in social networks. Signal Process. Image Commun. 57, 1–7.
8. Anmol, T., Sitara, K., 2024. Video source camera identification using fusion of texture
9. Apple Computer, I, 2001. Quicktime file format specification.
10. Bennabhaktula, G.S., Timmerman, D., Alegre, E., Azzopardi, G., 2022. Source camera device identification from videos. SN Computer Science 3, 316.
11. Bitmovin, 2023. Video developer survey 2023. URL: <https://bitmovin.com/video-dev>
12. eloper-report. (Accessed 15 July 2024).

BIOGRAPHIES

Nerella Bharath Ganesh studying 3rd year department of Electronics and Communication Engineering at Institute of Aeronautical Engineering (IARE) Dundigal. He published a research paper recently at ijsrem as a part of academics he has interest in IOT and Image Processing .



Dr Sonagiri China Venkateswarlu professor in the Department of Electronics and Communication Engineering at the Institute of Aeronautical Engineering (IARE). He holds a Ph.D. degree in Electronics and Communication Engineering with a specialization in Digital Speech Processing. He has more than 40 citations and paper publications across various publishing platforms, and expertise in teaching subjects such as microprocessors and microcontrollers , digital signal processing, digital image processing, and speech processing. With 20 years of teaching experience, he can be contacted at email: c.venkateswarlu@iare.ac.in



Dr. V. Siva Nagaraju is a professor in the Department of Electronics and Communication Engineering at the Institute of Aeronautical Engineering (IARE). He holds a Ph.D. degree in Electronics and Communication Engineering with a specialization in Microwave Engineering. With over 21 years of academic experience, Dr. Nagaraju is known for his expertise in teaching core electronics subjects and has contributed significantly to the academic and research community. He can be contacted at email: v.sivanagaraju@iare.ac.in.