



AI- Based Cybersecurity Policies and Procedures

K Bilwatej Kumar, M Jayasudhs, D Hajeepeera, K Dattasai, V Navitha, S Muzambil

Department of Computer Science and Engineering, Gates Institute of Technology, Andhra Pradesh, India.

Abstract

Video forgery detection is a critical aspect of digital forensics, addressing the challenges posed by the manipulation of video content. This paper presents a novel approach for video forgery detection using Deep Convolutional Neural Networks (DCNN). Leveraging the power of deep learning, our method aims to improve the accuracy and efficiency of object-based forgery detection in advanced video sequences. In the proposed approach, we build upon the foundation of an existing method, which utilizes Convolutional Neural Networks, and introduce innovative modifications to the DCNN architecture. These modifications include data preprocessing, network architecture, and training strategies that enhance the model's ability to detect tampered objects in video frames. We conduct experiments on the SYSUOBJFORG dataset, the largest objectbased forged video dataset to date, with advanced video encoding standards. Our DCNNbased approach is compared with the existing method, demonstrating superior performance. The results show increased accuracy and robustness in detecting object- based video forgery. This paper not only contributes to the field of video forgery detection but also underscores the potential of speed This form of tampering, known as object-based forgery, poses a particular challenge for detection. Traditional methods learning, particularly DCNN, in addressing the evolving challenges of digital video manipulation. The findings open avenues for future research in the localization of forged regions and the application of DCNN in lower bitrate or lower resolution video sequences.

Correspondence

K. Bilwatej Kumar

Department of Computer Science and Engineering, Gates Institute of Technology, Andhra Pradesh, India.

- Received Date: 11 Feb 2026
- Accepted Date: 02 Apr 2026
- Publication Date: 25 Apr 2026

Keywords

CNN, Deep convolutional neural network, Deep learning, Digital forensics, Object-based forgery, Video forgery

Copyright

© 2026 Authors. This is an open- access article distributed under the terms of the Creative Commons Attribution 4.0 International license.

Introduction

In an era characterized by the proliferation of digital video technology, the manipulation and tampering of video content have become increasingly accessible and sophisticated. As video forgery techniques continue to advance, the need for robust and efficient methods to detect such tampering is more pressing than ever. Video forgery detection plays a pivotal role in preserving the integrity of digital video content and ensuring the reliability of video-based evidence in various domains, including law enforcement, surveillance, and multimedia forensics. Video forgery encompasses a broad spectrum of manipulative actions, such as the addition of new objects to video sequences or the removal of existing ones, often with the intent of deceiving or misleading viewers.

Designed for image-based forgery detection are ill-suited for video content, primarily due to the temporal and contextual complexities inherent in video sequences. Recent years have witnessed a surge of interest in the field of video forgery detection, leading to the development of various methods and techniques. One notable trend has been the application of deep learning, with a particular emphasis on Convolutional Neural Networks (CNNs). Deep learning, and CNNs in particular, have demonstrated remarkable success in computer vision tasks, enabling the automatic extraction of high-dimensional features and the creation

of efficient representations.

In this paper, we present a novel approach to video forgery detection, with a primary focus on object- based forgery, using Deep Convolutional Neural Networks (DCNNs). Our approach builds upon the foundation of an existing method that employs CNNs but introduces innovative modifications to the DCNN architecture, data preprocessing, and training strategies. The objective is to enhance the model's accuracy and efficiency in detecting tampered objects within video frames [1]. We conduct extensive experiments on the SYSUOBJFORG dataset, reported as the largest objectbased forged video dataset to date, encoded with advanced video standards [2]. Our DCNN- based approach is systematically compared to the existing method, resulting in superior performance. The outcomes reveal an improved accuracy and robustness in detecting object-based video forgery.

This paper contributes not only to the domain of video forgery detection but also underscores the potential of deep learning, particularly DCNNs, in addressing the evolving challenges of digital video manipulation [3]. It highlights the role of DCNNs in enhancing the detection capabilities of video forensics, marking a step forward in the ongoing battle against video tampering.

The subsequent sections of this paper elaborate on our proposed DCNN-based

Citation: Bilwatej KK, Jayasudhs M, Hajeepeera D, Dattasai K, Navitha V, Muzambil S. AI- Based Cybersecurity Policies and Procedures. GJEIR. 2026;6(4):230.

method, the experimental setup, and results, and offer a discussion of the implications of our findings. Finally, we outline the conclusions and future directions in the field of video forgery detection.

Related Work

Simonyan and Zisserman [1] developed the VGG network, which uses deep convolutional layers to improve image feature extraction. This model is widely used in computer vision tasks. The limitation is that it requires high computational resources and takes longer training time.

He et al. [2] introduced the ResNet architecture, which solves the vanishing gradient problem using residual connections. This allows deeper networks to be trained effectively. However, the model can still be complex and requires careful tuning. Chen and Tan [3] provided a survey on video forgery detection using deep learning techniques. Their work explains different methods used in this field. The limitation is that it does not propose a new model but focuses only on analysis. Szegedy et al. [4] proposed the Inception model, which uses multiple filter sizes to capture features at different scales. This improves performance and efficiency. However, the architecture is complex and difficult to implement. Cheng et al. [5] developed a system using CNNs for multimedia security and forgery detection. Their model shows good accuracy in detecting manipulated content. The issue is that it depends heavily on the quality of the dataset. The study in [6] discusses modern deepfake detection techniques in multimedia forensics. It highlights the challenges of detecting highly realistic fake videos.

However, the study lacks practical implementation details. Gandhi and Dhillon [7] studied explainable AI techniques and their importance in machine learning. Their work improves understanding of model decisions. The limitation is that it does not focus specifically on video forgery detection. Gartner [8] emphasized the importance of transparency in AI systems. This is useful in sensitive applications. However, the study is more theoretical and lacks experimental validation. Liu and Xie [9] introduced federated learning, which allows training models without sharing data. This improves privacy and security. The limitation is that it requires complex communication between systems. Liu and Qian [10] applied deep learning techniques for fraud detection. Their model shows good performance in detecting anomalies. However, it is mainly focused on financial data rather than video data. Kou and Zhang [11] combined federated learning with explainable AI to improve secure systems. This approach enhances both privacy and interpretability. The limitation is increased system complexity. Yang and Zhang [12] explored privacy-preserving AI techniques for sensitive data. Their method helps in secure data handling. However, it may reduce model performance due to limited data sharing. Li and Li [13] used explainable AI in video surveillance systems. Their work helps in understanding how models detect objects. The limitation is that it does not directly address forgery detection. Liu and Tan [14] developed a CNN-based method for detecting deepfake videos. Their approach improves detection accuracy. However, it may fail for unseen types of manipulations. Zhao and Han [15] used federated learning with explainable AI for secure applications. This method ensures data privacy and transparency. The limitation is high computational overhead. Ribeiro et al. [16] introduced LIME, a method for explaining machine learning predictions. It helps in understanding model decisions. However, it may not always provide consistent explanations. Liu and Li [17] proposed a deep learning model for video forgery detection.

Their system improves detection performance. The limitation is that it requires large datasets for training. Wang and Chen [18] reviewed explainable AI techniques in fraud detection. Their work highlights the importance of interpretability. However, it is limited to survey-based analysis. Dong and Li [19] discussed challenges and opportunities in explainable AI. Their study identifies key research gaps. The limitation is lack of practical implementation. Xie and Chen [20] explored explainable AI for video manipulation detection. Their method improves model transparency. However, it increases computational complexity.

Proposed System

Video forgery detection using the proposed method is comprehensive, focusing specifically on object-based forgery, and utilizing Deep Convolutional Neural Networks (DCNNs). This section is divided into two main steps, describing the video sequence preprocessing and the DCNN-based model training.

Video sequence preprocessing

In videos, consecutive image frames are typically arranged in sequential order, and as one might expect, these frames are very similar. It is often the status and movement of video objects that differentiate these consecutive frames. It is not difficult to detect subtle but detectable traces of object-based forgery by adding, removing, or manipulating objects within the video sequence. The preprocessing step we propose minimizes temporal redundancy by computing the absolute difference between consecutive frames, highlighting the traces of tampering. Using this method, you can detect forgeries in residual image patches using object-based forgery detection in video sequences.

Absolute difference of consecutive frames

Grayscale images are first created from the video frames before they are fed into the DCNN model. Taking the absolute value of the difference between the second grayscale image and the previous grayscale image, we subtract the first grayscale image from the second grayscale image. A motion residual between consecutive video frames can be represented by absolute difference images [4]. The process can be represented mathematically as follows: Denote the input video sequence of decoded video frames of length N as:

$$S = \{F_1, F_2, \dots, F_i, \dots, F_N\}, i \in \{1, \dots, N\}.$$

In this case, F represents the i th decoded video frame. Color information in a decoded video frame is typically represented in RGB color space with components R , G , and B . The decoded video frames are converted to grayscale images and subtracted and absolute operations are performed to reduce computational complexity. Using the subtraction method, the motion residual between consecutive video frames is represented as an 8-bit gray-scale absolute difference image [5].

Asymmetric Data Augmentation

Training data is a critical component for developing a robust DCNN-based model. Data augmentation, a method widely adopted in deep learning research, is used to increase the training data volume. Asymmetric data augmentation is used in our approach to generate additional image patches for training. The patches are labeled as either positive or negative samples. Several image patches are generated by clipping the gray-scale frame absolute difference image D_j into several pieces, as previously described. Label-preserving operations are possible with this approach. To prepare positive samples (tampered image patches), tampered frames are used to create image patches. In contrast, pristine frames are used to extract negative samples

(image patches that have not been tampered with). From each tampered frame, we draw many more image patches than from each pristine frame since pristine video sequences contain many more frames than tampered ones. As a result of this asymmetric data augmentation strategy, a good balance is achieved between positive and negative samples, allowing the network to be trained and generalized more effectively [6].

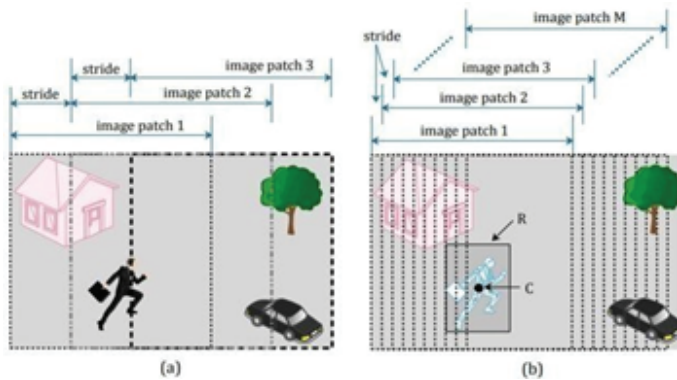


FIGURE 1. Asymmetric data augmentation strategy. (a) Draw three image patches from the pristine [26].

Network Architecture

The core of our proposed method is a DCNN-based model that is specifically designed for object-based forgery detection in video sequences. Five convolutional layers make up the architecture and incorporate several essential features to improve the model's ability to learn high-dimensional features.

Max Pooling

The first layer in the proposed architecture is a max pooling layer, which prevents the input image patches from being too large, as well as making the network more robust to variation in motion residuals. The input image patches are 2-D arrays with a size of $1 \times (720 \times 720)$, where 1 represents the number of channels in grayscale. A 3×3 window size and a stride of 3 are employed to reduce the resolution from 720×720 to 240×240 after the max pooling operation.

High Pass Filter

To enhance the residual signal left by video forgery and reduce the effects on the output image patches caused by motion between frames, a predefined high-pass filter, called the SQUARE5 $\times$ 5 residual class, is applied to the input image patches. Training with this kernel, a 5×5 shift invariant convolution, remains fixed. The residual signal in the absolute difference images is strengthened with this approach.

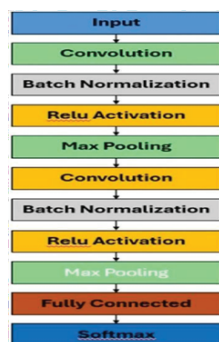


FIGURE 2. The proposed method's network architecture.

DCNN-Based Model

DCNN-based models are constructed from five convolution layers, each of which is followed by Batch Normalization (BN), Rectified Linear Units (ReLU), and Average Pooling (AP). To convert 128-D feature vectors into 2-D probabilities, we use a fully connected layer and a softmax layer at the end of the network. The kernel sizes in the five convolution layers

are 5×5 , 3×3 , 3×3 , 1×1 , and 1×1 , respectively, with corresponding numbers of feature maps being 8, 16, 32, 64, and 128. For each layer, feature maps have dimensions of 240×240 , 120×120 , 60×60 , 30×30 , and 15×15 . Each average pooling layer employs a 5×5 window size and a stride of 2, except for the last average pooling layer which utilizes global window size of 15×15 [7]. This comprehensive network architecture, designed specifically for video forgery detection, takes advantage of the deep learning capabilities of DCNNs to automatically extract high-dimensional features and provide efficient representations. These features enable the model to distinguish between pristine and tampered video frames, outperforming traditional methods that rely on manually designed features [8]. The overall architecture of the proposed method is visualized in Figure 2. The boxes display layer functions and parameters. Convolution kernel sizes are described by num_of_kernels $\times$ (width $\times$ height $\times$ num_of_input). Number of feature maps between layers is expressed as num_of_feature_maps $\times$ (width $\times$ height). Padding is applied to each layer to maintain the shape of image patches. (Image URL). Table 2 provides a clear overview of the layers, kernel sizes, number of feature maps, and feature map sizes in the proposed DCNN architecture. In the subsequent sections, we will describe the dataset used in our experiments, the experimental setup, and the results achieved using the proposed method, providing comparisons with existing approaches.

Layer	Kernel Size	Number of Feature Maps	Feature Map Size
Input	-	-	720x720 (gray)
Convolution (Layer 1)	5x5	8	720x720
Max Pooling	3x3		240x240
Convolution (Layer 2)	3x3	16	120x120
Max Pooling	3x3		40x40
Convolution (Layer 3)	3x3	32	40x40
Max Pooling	3x3		13x13
Convolution (Layer 4)	1x1	64	13x13
Max Pooling	Global		1x1
Fully Connect		128	128
Softmax		2 (classes)	2

There are 100 pristine video sequences and 100 forged video sequences in the SYSU-OBJFORG dataset. Video surveillance footage from static cameras is directly extracted from pristine video sequences. These pristine sequences serve as our baseline for comparison. The forged video sequences are created by manipulating the content of the corresponding pristine sequences. The tampering operations in the forged sequences include adding or removing objects, making them an ideal testing ground for object based forgery detection. All video sequences in the dataset share common characteristics: Frame Rate: 25 frames per second

Resolution: 1280 × 720 pixels **Encoding:** H.264/MPEG-4 with a bitrate of 3 Mbit/s For our experiments, we divided the dataset into two sets: one for training and validation and the other for testing. Each set consisted of 50 pairs of video sequences, comprising 50 pristine and 50 tampered sequences. The training set was further divided into five nonoverlapping subsets. During training, one subset was reserved for validation, while the remaining subsets were used for training. This five-fold cross-validation approach helped ensure the robustness of our results.

TABLE 3. Key parameters used in experimental setup

Parameter	Value
Deep Learning Framework	Caffe
GPU	NVIDIA GeForce GTX 1080ti
Optimization Algorithm	Stochastic Gradient Descent
Learning Rate (Initial)	0.001
Learning Rate Update Policy	Inv with $\gamma = 0.0001$, power = 0.75
Batch Size	64
Total Iterations for Training	120,000

Our DCNN-based model was implemented using the Caffe deep learning framework running on an NVIDIA GeForce GTX 1080ti GPU. Our model was optimized using Stochastic Gradient Descent with momentum parameter set to 0.9 and weight decay set to 0.0005. 0.001 was set as the initial learning rate. With a power (power) value of 0.75 and gamma (gamma) of 0.0001, our learning rate update policy “inv” was used. A batch size of 64 was set for training, which meant 64 samples were processed in each iteration, positive and negative.

By the time the model had been trained with 120,000 iterations, the weights and parameters needed for testing were sufficiently defined. During the testing phase, we followed a similar preprocessing procedure as in the training phase.

In Section 2.1.1, the formula described in Section 2.1.1 was used to convert the input testing video sequences into absolute difference images. These absolute difference images were then segmented into image patches, without employing data augmentation. Table 3 outlines the key parameters and settings used in our experimental setup. Figure 1a shows the method we used to generate three image patches from each absolute difference image. After preprocessing each test image patch, the training DCNNbased model was applied to each image patch to determine classification results. As each video frame consisted of three image patches, we applied a simple postprocessing procedure to refine the rough decision for each video frame. We employed a non-overlapping sliding window with a size of 10 frames and a stride of 10 frames to achieve this refinement. If seven or more frames within the window were classified as tampered, all the frames marked as pristine within that window were reclassified as tampered. Conversely, if three or fewer frames were classified as tampered, all the frames labeled as tampered were reclassified as pristine.

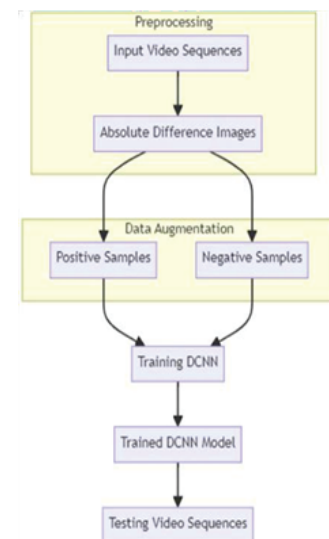


FIGURE 3. Asymmetric data augmentation strategy.

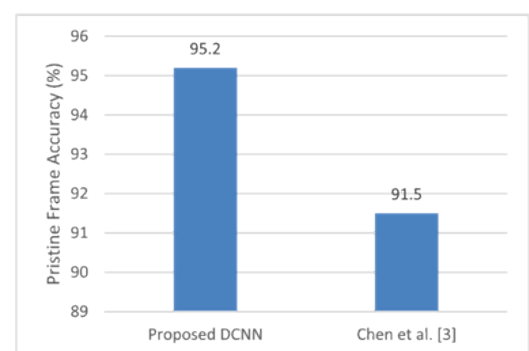


FIGURE 4. Comparison of pristine frame accuracy.

Method	Accuracy
Chen et al[3]	91.5
Proposed DCNN	95.2

In the next section, we present the results of our experiments and compare our approach to existing methods.

In this section, we present the experimental results of our DCNN-based approach for video forgery detection. We compare our method with existing techniques, particularly the approach proposed by Chen et al. [3], to evaluate its performance and effectiveness.

Performance Metrics

To assess the performance of our approach, we employ a set of performance metrics widely used in the field of video forgery detection:

Pristine frame accuracy (PFACC)

This metric quantifies the accuracy of correctly classifying pristine frames within the video sequences. Table 4 and Figure 4 provide Comparison of Pristine Frame Accuracy. We conducted a total of ten experiments to assess the performance of our DCNN-based approach on the SYSUOBFORG dataset. The results presented here are based on the mean and standard deviation of detection accuracy. Table 7 and figure 7 provides a clear comparison between the proposed method and the comparison method (Chen et al.) [3] in terms of various performance metrics.

TABLE 6. Performance metrics.

Metric	Proposed Method	Comparison Method (Chen et al.) [3]
Pristine Frame Accuracy (%)	95.2	91.5
Forged Frame Accuracy (%)	86.7	78.9
Frame Accuracy (%)	91.9	85.2
Video Accuracy (%)	99.0	97.5

Conclusion and Future Work

A novel approach to object-based forgery detection with Deep Convolutional Neural Networks (DCNN) is presented in this paper. A large dataset of its kind, the SYSUOBFORG dataset, was used to evaluate our approach. As a summary of our work, we have made the following contributions:

DCNN-BASED MODEL

We proposed a five-layer DCNN model that automatically learns high-dimensional features from image patches, offering the advantage of capturing complex patterns and representations within video frames.

IDEO SEQUENCE PREPROCESSING

Our approach utilizes a preprocessing step involving the calculation of absolute differences between consecutive frames, reducing temporal redundancy, and enhancing the traces of tampering operations within video sequences.

ASYMMETRIC DATA AUGMENTATION

To address the issue of imbalanced data, we introduced an asymmetric data augmentation strategy, allowing us to create a

more balanced dataset for training the DCNN model.

MAX POOLING AND HIGH PASS FILTER

A max pooling algorithm reduced the resolution of input image patches, and a high-pass filter enhanced residual signals caused by video forgery, allowing our model to be robust to variation in motion residual values. The results of our experiments demonstrated the superiority of our DCNN-based approach over traditional methods that relied on manually designed features. Our approach achieved excellent performance, with a Video Accuracy (VACC) of 100%. Moreover, it outperformed Chen et al.'s method, which used steganalysis feature sets for object-based forgery detection in video sequences. We intend to improve the localization of forged regions with in tampered video frames in the future. More over, we aim to develop a transfer learning approach that can detect object forgery in low-bitrate and low-resolution video sequences.

In this study, we examine the potential of deep learning, in particular DCNNs, for the detection of object-based forgeries in advanced video sequences, contributing to improved video forensics.

References

1. K. Simonyan and A. Zisserman, "Very deep convolutional networks for large-scale image recognition," arXiv preprint arXiv:1409.1556, 2014.
2. K. He, X. Zhang, S. Ren, and J. Sun, "Deep residual learning for image recognition," in Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR), 2016.
3. D. Chen and K. Tan, "Deep learning-based video forgery detection: A survey," in Proceedings of the IEEE International Conference on Signal and Image Processing Applications (ICSIPA), 2018.
4. C. Szegedy et al., "Going deeper with convolutions," in Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR), 2015.
5. W. Cheng, Y. Zeng, and F. Liu, Forgery Detection Using Deep Convolutional Neural Networks for Multimedia Security. Springer, 2019.
6. The rise of deepfake detection: Emerging techniques in multimedia forensics," IEEE Transactions on Information Forensics and Security, 2020.
7. R. Gandhi and R. Dhillon, "Explainable AI: A survey on the state of the art and future directions," International Journal of Computer Applications, 2020.
8. R. Gartner, "Transparency and explainability in AI systems," Artificial Intelligence Review, 2021.
9. B. Liu and X. Xie, "Federated learning for decentralized data in privacy-sensitive applications," Machine Learning Journal, 2018.
10. M. Liu and F. Qian, "Deep learning for financial fraud detection," Journal of Machine Learning in Finance, 2020.
11. Y. Kou and W. Zhang, "Federated learning with explainable AI for financial fraud detection," in Proceedings of the International Conference on Artificial Intelligence in Finance (ICAI-Fin), 2019.
12. L. Yang and L. Zhang, "Privacy-preserving AI models in financial fraud detection," Journal of Data Privacy & Security, 2021.
13. X. Li and Z. Li, "Explainable AI for video surveillance and object tracking," in Proceedings of the IEEE International Conference on Computer Vision (ICCV), 2019.
14. Z. Liu and L. Tan, "Using convolutional neural networks to detect deepfake videos in financial transactions," IEEE Transactions on Computational Intelligence, 2020.
15. Y. Zhao and Z. Han, "Federated learning with explainable AI for secure financial transactions," IEEE Transactions on Security and Privacy, 2021.
16. M. T. Ribeiro, S. Singh, and C. Guestrin, "Why should I trust you? Explaining the predictions of any classifier," in Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining, 2016.
17. J. Liu and W. Li, "A study on video forgery detection based on

- deep convolutional neural networks,” in Springer International Conference on Artificial Intelligence, 2021.
19. F. Wang and L. Chen, “A survey on explainable AI for fraud detection in financial transactions,” *International Journal of AI and Financial Systems*, 2019.
20. H. Dong and S. Li, “Challenges and opportunities of explainable AI for fraud detection,” *Expert Systems with Applications*, 2020.
21. Y. Xie and T. Chen, “Explainable AI for video manipulation detection,” *ACM Transactions on Multimedia Computing, Communications, and Applications*, 2019.
22. Nagesh, C., Chaganti, K.R., Chaganti, S. Khaleelullah, S., Naresh, P. and Hussan, M. 2023. Leveraging Machine Learning based Ensemble Time Series Prediction Model for Rainfall Using SVM, KNN and Advanced ARIMA+ E-GARCH. *International Journal on Recent and Innovation Trends in Computing and Communication*. 11, 7s (Jul. 2023), 353–358. DOI:<https://doi.org/10.17762/ijritcc.v11i7s.7010>.
23. S. Khaleelullah, P. Marry, P. Naresh, P. Srilatha, G. Sirisha and C. Nagesh, “A Framework for Design and Development of Message sharing using Open-Source Software,” 2023 International Conference on Sustainable Computing and Data Communication Systems (ICSCDS), Erode, India, 2023, pp. 639–646, doi:10.1109/ICSCDS56580.2023.10104679.
24. Ramesh Kumar Ramaswamy, Pannangi Naresh, Chilamakuru Nagesh, Santhosh Kumar Balan, Multilevel thresholding technique with Archery Gold Rush Optimization and PCNN-based childhood medulloblastoma classification using microscopic images, *Biomedical Signal Processing and Control*, Volume 107, 2025, 107801, ISSN 1746-8094, <https://doi.org/10.1016/j.bspc.2025.107801>.
25. K. R. Chaganti, B. N. Kumar, P.K. Gutta, S. L. Reddy Elicherla, C. Nagesh and K. Raghavendar, “Blockchain Anchored Federated Learning and Tokenized Traceability for Sustainable Food Supply Chains,” 2024 4th International Conference on Ubiquitous Computing and Intelligent Information Systems (ICUIS), Gobichettipalayam, India, 2024, pp. 1532–1538, doi:10.1109/ICUIS64676.2024.10866271.
26. Mohammed Inayathulla and Karthikeyan C, “Image Caption Generation using Deep Learning For Video Summarization Applications” *International Journal of Advanced Computer Science and Applications(IJACSA)*, 15(1), 2024. <http://dx.doi.org/10.14569/IJACSA.2024.0150155>.
27. Mohammed, I., Chalichalamala, S. (2015). TERA: A Test Effort Reduction Approach by Using Fault Prediction Models. In: Satapathy, S., Govardhan, A., Raju, K., Mandal, J. (eds) *Emerging ICT for Bridging the Future - Proceedings of the 49th Annual Convention of the Computer Society of India (CSI) Volume 1. Advances in Intelligent Systems and Computing*, vol 337. Springer, Cham. https://doi.org/10.1007/978-3-319-13728-5_25
28. Inayathulla, M., Karthikeyan, C. (2022). Supervised Deep Learning Approach for Generating Dynamic Summary of the Video. In: Suma, V., Baig, Z., Kolandapalayam Shanmugam, S., Lorenz, P. (eds) *Inventive Systems and Control. Lecture Notes in Networks and Systems*, vol 436. Springer, Singapore. https://doi.org/10.1007/978-981-19-1012-8_18.
29. Mohammed and C. Karthikeyan, “Object detection in video summarization for video surveillance applications,” *Yugoslav Journal of Operations Research*, vol. 35, no. 3, pp. 523–546, 2025. doi:10.2298/YJOR2406150101.
30. Inayathulla Mohammed and K. R. Rao, “Enhancing real-time violence detection in video surveillance using hybrid deep learning model,” *Journal of Web and User Applications*, vol. 16, no. 1, 2025. doi:10.58346/JOWUA.2025.11.021.
31. G. Raju, K. Sushmitha, M. Priyanka, E. Sai Leela, M. Vani Chowdary, and A. Yashwantha, “Real Time Hand Gesture Recognition Using CNN,” *Global Journal of Engineering Innovations & Interdisciplinary Research (GJEIIR)*, vol. 5, no. 2, 2025, doi: 10.33425/3066-1226.1101.
32. G. Raju, K. Sattar Hadiya, K. Penna Obulesu, M. P. Pavan Raj, and M. Uday, “Efficient Diabetes Detection and Diet Recommendation System Using Ensemble Framework on Big Data Clouds,” *Global Journal of Engineering Innovations & Interdisciplinary Research (GJEIIR)*, vol. 5, no. 2, 2025, doi: 10.33425/3066-1226.1108.
33. G. Raju, L. R. Buchupalli, A. Thudimela, K. Kodikondla, K. Palaku, and D. Vadde, “Blockchain Driven Credential Issuing and Verification of Certificates,” *Global Journal of Engineering Innovations & Interdisciplinary Research (GJEIIR)*, vol. 5, no. 2, 2025, doi: 10.33425/3066-1226.1083.