

Original Article

DCA: A Residual CNN And Wireshark-Based Framework for Detecting Double JPEG Compression Anti-Forensic Attacks in Cloud Environments

Shaik Sharmila¹, Ch. Aparna²

¹Department of Computer Science and Engineering, Acharya Nagarjuna University, Guntur, Andhra Pradesh, India.

²Department of Computer Science & Engineering, R.V.R & J.C. College of Engineering, Guntur, Andhra Pradesh, India.

¹Corresponding Author : mail2shaiksharmila@gmail.com

Received: 20 March 2026

Revised: 21 April 2026

Accepted: 17 June 2026

Published: 29 July 2026

Abstract - Cloud computing brings integrity and reliability of digital images with a spectrogram of an anti-forensics method named as the double JPEG compression. These are meant to conceal the traces of intrusion which poses difficulties to the forensic analysts and investigators. In this direction, we come up with DCA (Detection and Classification Architecture), a hybrid and strong framework that does attack-resistant detection and localization of anti-forensic of the double JPEG compression. Network architecture incorporates the best capabilities of CNNs in extracting features and the Wireshark-based network traffic analysis, which provides improved traceability and validation. The CNN model can extract spatial and frequency domain inefficiency in image data that cannot be seen by a human eye but rather are recompression artifacts. Unlike conventional methods that rely solely on image-domain analysis, the proposed DCA introduces a dual-domain forensic strategy that integrates spatial-frequency feature learning with network-level traffic monitoring, enabling both content-based and transmission-aware verification of image authenticity. Meanwhile, it monitors transmission patterns in Wireshark to monitor the spread of images that were manipulated, and where it all goes amiss. The outputs of the system include binary classification (tampered/untampered) and heatmaps to overlay on the input images to show the suspicious areas. To test the quality of DCA, we tested DCA on emphasized performance measures: Precision, Accuracy, True Positive Rate (TPR), False Positive Rate (FPR), F-score, Peak Signal-to-Noise Ratio (PSNR), and Structural Similarity Index Measure (SSIM). Experimental evidence indicates that DCA is interpretably stable with high detectability. The main input of the work is the creation of a coherent multi-domain forensic model that combines the residual CNN-based spatial frequency feature extraction with network traffic analysis Wireshark. The proposed DCA framework can detect and trace multiple tampering of images simultaneously compared to the conventional techniques that use single-domain analysis, leading to better accuracy, robustness, and reliability when using same-quantization double JPEG compression configurations. The framework also offers classification as well as localization results and is best suited to the real-time cloud forensic investigations.

Keywords - Anti-forensic detection, Cloud forensics, Convolutional Neural Network, Double JPEG compression, Image tampering.

1. Introduction

With the fast development of digital imaging and Artificial Intelligence technology, the occurrence of advanced image manipulation and forgery attacks has been boosted in modern multimedia platforms. With the advent of AI-powered image synthesis systems, sophisticated editing software, and generative models, it is now possible to produce more realistic-looking counterfeit images, especially as they become hard to spot with traditional image forensic methods. As a consequence, digital image authentication has been an important research field in various areas, including cyber forensics, healthcare systems, legal investigations, journalism,

and intelligent surveillance applications. The ability to manipulate images and the number of images produced by AI have increased rapidly, giving digital image forensics a lot of attention. One of the difficult issues in image forensics is the detection of double JPEG compression with the same quantization matrix proposed by Niu et al. [1].

To detect AI-generated images, Yu and Xu [2] designed a multi-modal texture fusion network, leveraging complementary visual features. Kadha et al. [3] presented an exhaustive survey of image manipulation detection and localization techniques and emphasized recent developments



of forensic systems based on Deep Learning. Moreover, Fengyong Li et al. [4] proposed a deep residual learning artifact removing network for JPEG images, which proved the impact of image enhancement techniques on forensic traces. These studies highlight the need to establish strong forensic strategies that can address compression artifacts, image manipulations, and synthetic content creation. In recent years, deep learning approaches have been increasingly used to enhance the accuracy and robustness of detection forgery.

A network named RED-Net was proposed by Chen et al. [5] for image steganalysis and by Aryan et al. [6] for efficient forgery detection in images, which is also a lightweight network using a self-attention mechanism. Rao et al. [7] presented a JPEG-resistant forgery detection framework based on self-supervised domain adaptation to improve the performance under different compression levels. In the same train of thought, Wei et al. [8] introduced JPEG-compression related artifacts and multi-dilated feature fusion into JAMD-Net for image splicing detection, while Zhao et al. [9] developed multi-scale feature fusion for document image tampering localization.

In summary, these methods reveal the strength of attention-based architectures, fusion of the features, and learning for compression are useful to improve modern digital image forensic systems. Although digital image forensics has made significant progress, the existing image forensic techniques still face difficulties in achieving lightweight deployment with strong compression resistance and high precision, as well as strong adaptability for different types of real-world images. Hence, there is a need for building a smart, adaptive, and lightweight image forensic framework which can effectively detect advanced image manipulations in the tough multimedia adversarial environment.

Despite these advancements, existing approaches predominantly focus either on image content analysis or network-level monitoring in isolation, lacking a unified framework that jointly leverages both domains for enhanced anti-forensic detection. Furthermore, many existing methods struggle under the same-quantization double compression and sophisticated anti-forensic conditions, indicating the necessity for more robust and integrated detection strategies. To address these limitations, this work proposes a hybrid Detection and Classification Architecture (DCA) that integrates residual CNN-based spatial-frequency feature extraction with network-level traffic analysis, enabling a multi-domain forensic validation mechanism for improved detection reliability in cloud environments.

The rationale of this study on the anti-forensic detection of a double JPEG compression in a cloud platform is informed by the growing adoption of cloud-based image sharing and image storage platforms, where recompression is often used. The presence of such operations can leave evidence of such

tampering unseen, and thus it can be extremely difficult to detect, especially when conditions of the same quantization are met. Current methods of forensic analysis have drawbacks in the manner that they can identify such manipulations. Thus, the combination of the analysis of the image domain with network-level surveillance, with the help of Wireshark, is a reliable and efficient tool that will enhance the forensic traceability and detection rates.

2. Literature Review

Anti-forensics has grown considerably, and its goal is to disguise any signs of image tampering and to frustrate the forensic system in determining the authenticity of the image. New works have involved more advanced methods which interfere with the localization of forgery and compression artifact detection. Zhuo et al. [10] designed an improved anti-forensic scheme, which can actively deceive the forgery localization scheme by altering the forensic traces without impacting visual quality.

In previous works, Sutthiwan et al. [11] were studying on anti-forensic approaches to hide the evidence of the JPEG double compression and succeeded to show the compression artifacts can be manipulated to prevent the detection of JPEG double compression. In a similar vein, Li et al. [12] also introduced a targeted anti-forensic technique for double JPEG compression in which the same quantization matrix is used and showed that the efficiency of forensic detectors can seriously be jeopardized by carefully manipulating compression traces.

In recent times, NM Uddin et al. [13] used adversarial generative networks to create visually realistic images to evade double JPEG compression detectors, further emphasizing the importance of deep learning in anti-forensic research. In total, these studies show that anti-forensic techniques are increasingly becoming sophisticated and adaptive, making them a significant threat to traditional forensic systems.

The combination of Generative Models, Statistical manipulation Techniques, and Adversarial learning mechanisms have contributed to enabling attackers to hide image processing histories. Hence, it has become a potentially important research direction in digital image forensics to build up strong forensic schemes that can withstand such sophisticated anti-forensic attacks.

One of the best studied issues in digital image forensics is double JPEG compression, which compression inconsistencies are seen to be very useful in image tamper detection. Weixuan Tang et al. [14] proposed a Deep Reinforcement Learning (DRL) anti-forensic framework that can learn an optimal strategy to mask the double compression processing artifacts, which proves the feasibility of using intelligent agents in anti-forensic scenarios.

On the contrary, Pevný and Fridrich [15] created one of the basic techniques for detection double JPEG compression, and it is based on the statistical anomalies created during the double JPEG compression. Based on the traditional statistical methods, Barni et al. [16] used a method based on Convolutional Neural Networks to detect aligned and non-aligned double JPEG compression, showing greater robustness against various compression cases.

On the other hand, Yang et al. [17] presented a detection scheme based on intra-DCT correlation, which is more effective in detecting double compression traces. To achieve better detection, B Li et al. [18] devised a multi-branch CNN, which is specifically tailored to double JPEG compression analysis. All these studies are witness to an ongoing improvement in methods for detecting forensic features, from handcrafted statistical features to advanced Deep Learning Architectures.

But the growth of anti-forensic countermeasures has resulted in a continuous cat-and-mouse game between forensic analysts and attackers. Hence, future forensic systems need to include the approaches of adaptive learning, explainable decision making, and robust feature extraction to effectively address the new anti-forensic challenges and provide a high detection rate under various image manipulation cases.

F Ding et al. [19] came up with ExS-GAN, which is an enhanced form of a GAN, which includes additional supervision that allows the network to achieve better anti-forensic performance in image editing tasks, but the research fails to state the dataset environment and a trade-off between the invisibility and quality of an image.

D Kim et al. [20] proposed an end-to-end CNN-based design aimed at resisting JPEG and Double JPEG (DJPEG) anti-forensic, which combines the composition of EDSR and DCT constraints and composite losses, which proves effective on resisting JPEG-based forensic detection with different datasets like Bossbase 1.01 and BOWS2, yet it suffers performance loss in non-aligned cases, and previous methods are time-consuming.

H Zhang et al. [21] proposed a framework based on GAN with a double-mask guided local perturbation algorithm that can evade face forensic CNN detectors, namely applied to anti-forensics in the face domain of GAN-generated face images, but the algorithm is face-specific and needs to be improved in terms of cross-model transferability.

In an anti-forensic context, X Feng et al. [22] developed a GAN-based which uses Residual Attention Fusion and chroma difference loss (to remove JPEG decompression artifacts), which can operate on JPEG decompressed image datasets (not explicitly stated), although the method does not fully utilize all

available input image feature information. Despite these developments, currently, the available methods either operate on image-domain techniques or network-level techniques in isolation without a unified framework that interactively uses the two techniques to achieve better anti-forensic detection.

Moreover, most of the available techniques fail in same-quantization double compression and advanced anti-forensics, leading to a lower level of accuracy and resilience in detection. Moreover, the lack of integrated feature learning at the spatial, frequency, and transmission-level data impede the reliability of the existing forensic systems in dynamic cloud environments. Thus, the main issue this paper aims to solve is to create a shared multi-domain forensic solution that can effectively identify anti-forensic attacks in the form of the doubling of the JPEG compression, as well as be robust, scaled, and enhanced traceability in the cloud environment.

3. Proposed Approach

The proposed DCA architecture, as shown in Figure 1, presents a hybrid, multi-stage pipeline for detecting anti-forensic manipulations, specifically double JPEG compression attacks, in cloud environments. It begins with data retrieval from a cloud-based database, which houses JPEG-compressed images.

These images are processed through a dual path: one going through Wireshark for network-layer analysis and the other undergoing feature extraction and data segmentation. Extracted features are stored in a centralized data storage module and then pre-processed and fed into a Hybrid CNN. This Residual-Infused Convolutional Neural Network includes skip connections and DCT layers that amplify subtle discrepancies introduced during double compression.

At the same time, the Wireshark stream is forwarded to another stream of the same CNN that symbolizes the integration of the evidence on the packet level. The processed data undergoes the feature selection process, where the pertinent input signals are improved, and the noise is removed. After some deep feature learning, the model carries out classification, which determines whether tampering has been made or not. In the event of the tampering being detected, the system identifies it as copy-move forgery, and a moral reprocessing is carried with a binary mask and morphological processing, used to identify the manipulated areas.

Measures of performance such as TPR, PSNR, and SSIM are calculated to measure the integrity of the output. The signal marking the detection is optionally fed back used to update or refine the feature selection or preprocessing phase to repeat the learning or forensic log.

This architecture is closely coupled with cloud forensics settings with the CGI-based UI front ends and allows conducting real-time forensic analyses, report creation, and

visualization. Wireshark has been included to make sure that the packet-level monitoring is robust, and the hybrid CNN approach is used to perform the spatial and frequency domain analysis of deep forensic validation.

Residual-Infused CNN for Anti-Forensic Attack Detection in Double JPEG Compression is a CNN-based method to identify tampered regions in double-compressed

JPEG images, which is in demand in the field of digital image forensics. The input JPEG image is first loaded by the algorithm, that converts it to a grayscale image for computational purposes and considering only the main structures. This is done so by computing a weighted average across channels that keeps important information and gets rid of some redundant data. We normalize the grayscale image to the range [0, 1] to obtain a stable pixel intensity value as learning from a Deep Learning Model.

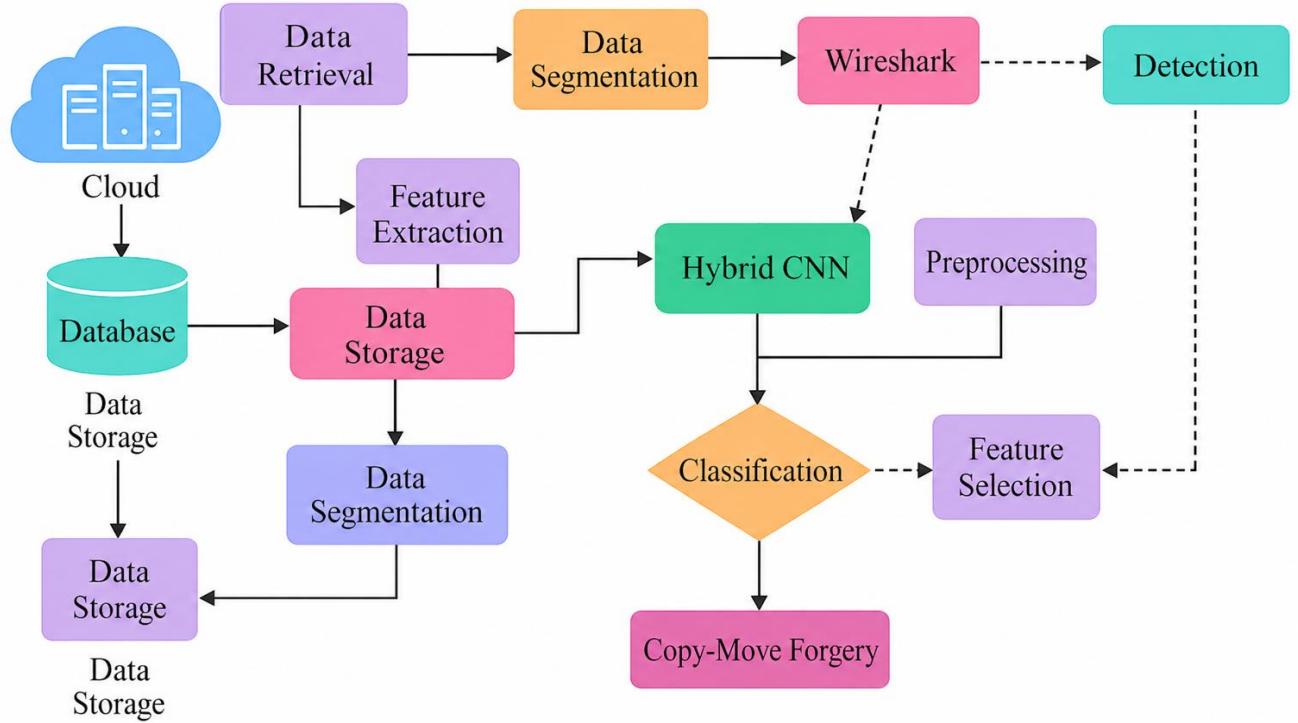


Fig. 1 Proposed architecture for detection copy-move forgery attacks

4. Proposed Approach to Algorithm

Algorithm 1 Residual-Infused CNN for Detecting Anti-Forensic Attacks in Double JPEG Compression

Input: JPEG compressed image I_{jpeg}

Output: Detection of anti-forensic manipulations and segmented output F

Step 1: Load image: Load the JPEG compressed image $I_{jpeg} \in \mathbb{R}^{m \times n \times 3}$.

Step 2: Convert to grayscale: $I_{gray} = \sum_{c=1}^3 W_c \cdot I_{jpeg}^{(c)}$ where W_c are channel weights.

Step 3: Normalize pixels: $I_{norm} = \frac{I_{gray}}{255}$ normalize pixel values to the range [0,1].

Step 4: Define convolution: Set up the convolutional filter W_{conv1} and bias b_{conv1} for feature extraction.

Step 5: Apply convolution: $f_{conv1} = W_{conv1} * I_{norm} + b_{conv1}$ convolve the normalized image I_{norm} .

Step 6: Apply ReLU: $f_{conv1}^{relu} = \max(0, f_{conv1})$ apply ReLU activation.

Step 7: Residual block: $f_{res1} = W_{res1} * f_{conv1}^{relu} + b_{res1}$ set up residual block filters.

Step 8: Skip connection: $f_{res1}^{out} = \max(0, f_{res1} + f_{conv1}^{relu})$ add skip connection.

$\Delta f_{res} > \epsilon f_{res2} = \max(0, W_{res2} * f_{res1}^{out} + b_{res2} + f_{res1}^{out})$

Step 9: perform residual learning. $\Delta f_{res} = |f_{res2} - f_{res1}|$ compute residual difference.

Step 10: Batch normalization: $f_{bnorm} = \frac{f_{res2} - \mu_{batch}}{\sigma_{batch} + c}$ where μ_{batch} and σ_{batch} are the batch mean and variance.

Step 11: Extract high-level features: $f_{high} = W_{high} * f_{bnorm} + b_{high}$ apply convolutional layer, ReLU activation: $f_{high}^{relu} = \max(0, f_{high})$.

Step 12: DCT on features: $C_{dct}(u, v) = \sum_{x=0}^{N-1} \sum_{y=0}^{N-1} f_{high}(x, y) \cos\left(\frac{(2x+1)u\pi}{2N}\right) \cos\left(\frac{(2y+1)v\pi}{2N}\right)$ apply 2D DCT.

Step 13: Extract DCT coefficients: Extract high-frequency DCT coefficients from $C_{dct}(u, v)$.

Step 14: Residual differences: $\Delta f_{res} = |f_{high} - f_{res2}|$ compute residual difference between DCT and high-level features. $\Delta f_{res} > \tau AF_{detected} = 1$ classify as tampered. $AF_{detected} = 0$ classify as untampered.

Step 15: Generate mask: $B[i] = \begin{cases} 0, & \text{if tampered} \\ 1, & \text{if untampered} \end{cases}$ create binary mask.

Step 16: Post-processing: Apply morphological operations on the binary mask B .

Step 17: Remove false positives: Remove small, connected components to reduce false positives.

Step 18: Classification: $y_{out} = \frac{\exp(W_{fc} \cdot AF_{detected} + b_{fc})}{\sum_j \exp(W_{fc} \cdot AF_{detected} + b_{fc})}$

Classify the tampered regions using a fully connected layer.

Step 19: Compute loss: $\mathcal{L}_{CE} = -\sum_i y_i \log(y_{out})$ compute cross-entropy loss for classification.

Step 20: L2 regularization: $\mathcal{L}_{total} = \mathcal{L}_{CE} + \lambda_1 \|W_{res1}\|_2^2 + \lambda_2 \|W_{fc}\|_2^2$ add L2 regularization to avoid overfitting.

Step 21: Backpropagation: $W^{(t+1)} = W^{(t)} - \eta \cdot \frac{\partial \mathcal{L}_{total}}{\partial W}$ Update weights using gradient descent.

Step 22: Evaluate performance: Compute performance metrics. True Positive Rate (TPR): $TPR = \frac{TP}{TP+FN}$ and Peak

Signal-to-Noise Ratio (PSNR): $PSNR = 10 \log_{10} \left(\frac{MAX^2}{MSE} \right)$.

Step 23: Post-processing: Generate the final segmented image F by applying the binary mask.

Step 24: Output: The resulting segmented image F with the tampered regions clearly marked.

Following the preprocessing of the input image, a first convolution layer is employed to generate a set of low-level features. This layer retains some low-level texture and edge information useful for identifying the compression artifacts. The first features are then passed through several residual blocks, which make up the core of the model. These blocks are equipped with skip connections to shuttle the fine-grained local information from low layers to high layers and hence to make complex patterns in late layers efficiently.

This structure reduces the gradient vanishing problem and allows the network to learn the subtle anti-forensic traces that are difficult to capture by the classic convolutional networks. The feature maps are normalized using batch normalization to stabilize training after the residual blocks. It follows a high-level convolutional layer connected to the non-linear activation function (ReLU) to capture a more abstracted structure. To further detect anti-forensic operations, the

algorithm applies a 2D Discrete Cosine Transform (DCT) to produce high-frequency contents (which may expose the forgery). The residual difference of the high-level features with the previous residual layer is computed to highlight differences, which became prominent due to the manipulation.

Then a binary mask is produced to distinguish between the tampered and untampered pixels by applying a threshold to the residual differences. This mask is post-processed using morphological operations to wipe off noise and decrease false positives so that the segmentation of tampered regions becomes neater.

The last layer is a dense layer for the final prediction with cross-entropy loss and L2 regularization to ensure the generality of the model and prevent over-fitting. Subsequently, the model is evaluated in terms of performance, for example, TPR and PSNR, in detecting anti-forensic attacks. The segmented image is an output of the algorithm, with the image showing distinguished tampered regions. The method has strong and accurate tampering detection for digital forensics.

5. Experimental Setup

A computer consisting of an i7 3.7GHz processor with 16GB RAM, 1TB NVMe storage is employed to implement the project of this FCSRNN. Microsoft Visual Studio IDE is used for developing a diligent UI (User Interface) that can load the algorithms and run them on a rented cloud server from i2k2. com. The CGI is supported, which means portability, as the scripts can be written in any programming language. It is a simple way for web servers to execute programs and produce output dynamically by handling interactivity and input from the user. CGI is easy to use, with no complex setup necessary, and can be used for smaller projects. And it rocks in low server overhead, or low dynamic content requirements. A versatile and powerful tool for extracting and tampering with metadata in image files is employed as a service in image forensic examination.

One of the main limitations is the fact that ExifTool can be used very straightforwardly to tamper with metadata, with no clear signs that any editing was performed, which is why authenticating content based only on the metadata is challenging. It does not have strong tampering detection or a deeper analysis of image integrity, like Error Level Analysis (ELA) or pixel-level inconsistencies. The tool also does not provide a user-friendly environment for analysis, which might cause problems to novice users in conducting forensic analysis. So it's also restricted to just getting metadata.

The dedicated UI, CGI, leased server, and the ExifTool Service are provided as the clean testbed, including the examination of the DCA work. To ensure reproducibility of the proposed framework, all experimental modules are

organized in a modular pipeline where preprocessing, feature extraction, CNN-based classification, and post-processing stages are independently executable. The implementation parameters, including normalization schemes, image resizing protocols, and feature extraction settings, are maintained consistently across all experiments. Furthermore, standardized input formats (JPEG images with defined compression parameters) are used to guarantee consistency in evaluation.

The system also logs intermediate outputs such as feature maps, DCT coefficients, and classification results, enabling step-by-step validation of the forensic pipeline. The integration between the UI, CGI interface, and backend processing modules is designed to allow repeatable execution under identical conditions, thereby supporting transparency and reproducibility in cloud-based forensic analysis.

The DCA interface, as reflected in Figure 2, is a clean and structured Graphical User Interface (GUI) designed to facilitate the anti-forensic attack detection process through a modular and user-friendly workflow. It allows users to set the working directory, upload datasets (e.g., FDS.zip), and configure essential components such as the gateway executable (CCSG.exe), Forensic Tool (FET.exe), and report file path. With dedicated buttons for launching CGI services, Internal Forensic Tools (IFT), loading datasets, performing analysis, generating reports, and visualizing results through graphs, the interface ensures streamlined execution of all cores forensic tasks. The status bar at the bottom clearly communicates successful operations, such as "Report generated successfully," enhancing usability for both technical and non-technical users involved in digital forensic investigations.

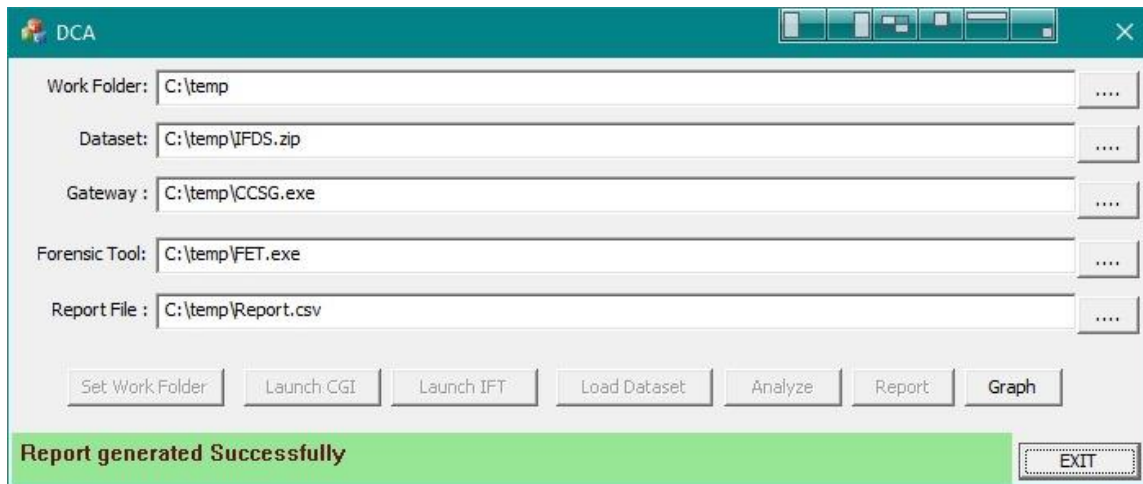


Fig. 2 DCA interface for digital forensic analysis workflow and report generation

6. Results and Discussions

The evaluation of the performance of the proposed work had been done in a comprehensive manner with various evaluation measures. The classification measures of accuracy, precision, TPR, and false positive rate were used in ratings of the model performance in its true predictions and false predictions.

All these metrics shed some light at the degree to which the system is operating with respect to repeatability. Additionally, we have also examined the F1-score to have a trade-off combination of the precision and recall of the model, especially when dealing with class imbalance. This will ensure that the model is consistent in terms of true positive identification and reduction of false negatives.

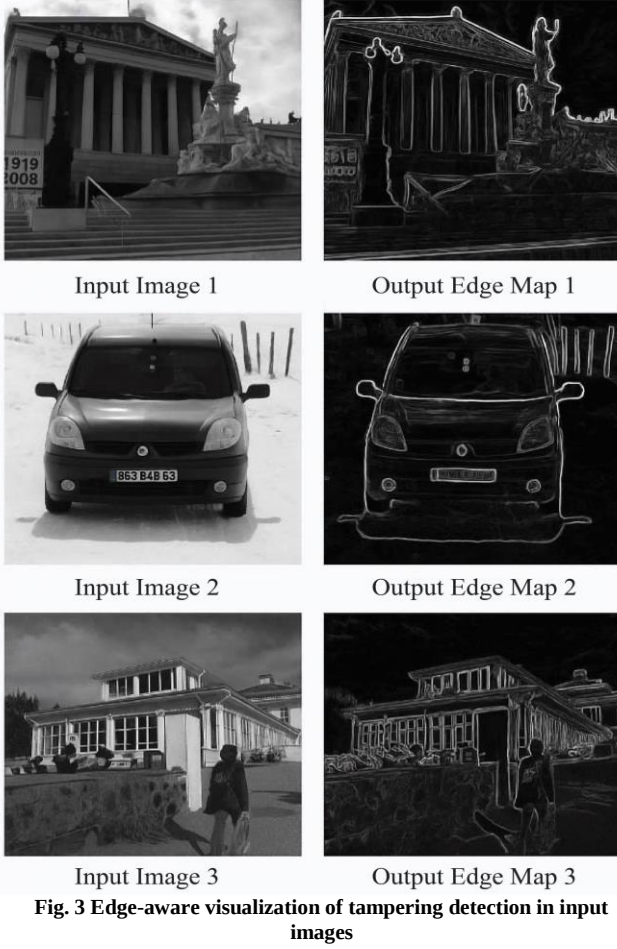
When using image quality assessment metrics, PSNR and SSIM were used as measures of visual quality and perceptual similarity of the source and reconstructed images. These factors contribute to the robustness of the model in preserving

image quality and structure, which is crucial for image-based applications.

Figure 3 below reflects a 3×2 grid that displays the sample input and output results of the proposed DCA framework. The left column contains the original grayscale input images (Input Image 1 to 3), while the right column shows the corresponding edge maps (Output Edge Map 1 to 3) generated by the Residual-Infused CNN. These edge maps highlight high-frequency regions and structural inconsistencies often linked to double JPEG compression or anti-forensic manipulations.

6.1. Sample Input and Output

It is noteworthy that such finer details are also preserved and stressed in the edge outputs like building contours, cars, or human silhouettes, which is of great help in the forensic analysis. This semiconservative evidence proves the capability of the model to localize the tampered areas by identifying the slight compression artifacts, which are not visible to the human eye, thus supporting the use of the model in real-time cloud forensics.



6.1.1.1. Accuracy

Accuracy is used to measure the general correctness of a classification model. It is the proportion of predictive correct observations (positives and negatives) to the aggregate number of observations. The relative accuracy between the different anti-forensic detection techniques represented in Table 1, such as EXSGAN, E2ENSJD, LPGMGFA, and JDGAN, and the proposed DCA framework indicates the high

level of accuracy with DCA. At the entire scope of the evaluation data (3 to 30 percent), DCA is always best and has the highest accuracy values up to 99.64088 percent. In contrast, the other methods, such as EXSGAN and JDGAN, show relatively lower accuracy, with values mostly fluctuating around 97.1% to 97.6%, indicating less robustness and generalizability.

LPGMGFA and E2ENSJD achieve an intermediate level of accuracy, reaching even up to 98.9%, but they lag behind DCA. The relatively small variance of performance of DCA among multiple percentages of available data (close to [99.49382%, 99.64088%]) indicates that the proposed model is solid and accurate to detect the double JPEG anti-forensic attacks in cloud computing. These results also justify that DCA can still obtain high prediction accuracy under different data distributions, making it a promising method for cloud-based digital forensic analysis.

The graph shown in Figure 4 reflects the relative performance of the five models based on the effectiveness of classification of different data sizes. As can be seen, the DCA has been scoring the highest performance values, which means that it is better at retaining both the precision and recall of all data points. LPGMGFA comes next, indicating stable and high performance with deviations of minimal variations as the data size advances.

E2ENSJD portrays a medium performance with a steadily high performance, but marginally below LPGMGFA. JDGAN demonstrates a relatively lower performance, and significant discrepancies in performance could be observed compared to the best models. EXSGAN has the least values of performance, meaning that it has some constraints on how to manage the data effectively. Overall, the graph shows that all the models have a stable tendency as the data increases; however, DCA is obviously the best among them, which is why it can be taken as the most credible method due to its overall classification performance.

Table 1. Comparative accuracy analysis (%) of anti-forensic detection models across varying data percentages

Accuracy (%)					
Data (%)	EXSGAN	E2ENSJD	LPGMGFA	JDGAN	DCA
3	97.32588	98.63705	98.80206	97.50206	99.55264
6	97.26706	98.51941	98.94913	97.50206	99.58206
9	97.17883	98.49	98.83147	97.64912	99.49382
12	97.14941	98.63705	98.8903	97.47265	99.49382
15	97.14941	98.54882	98.86089	97.50206	99.64088
18	97.26706	98.57823	98.83147	97.47265	99.58206
21	97.3553	98.54882	98.83147	97.44324	99.61147
24	97.26706	98.60764	98.71382	97.61971	99.61147
27	97.12	98.7253	98.74324	97.53147	99.55264
30	97.32588	98.57823	98.71382	97.47265	99.49382

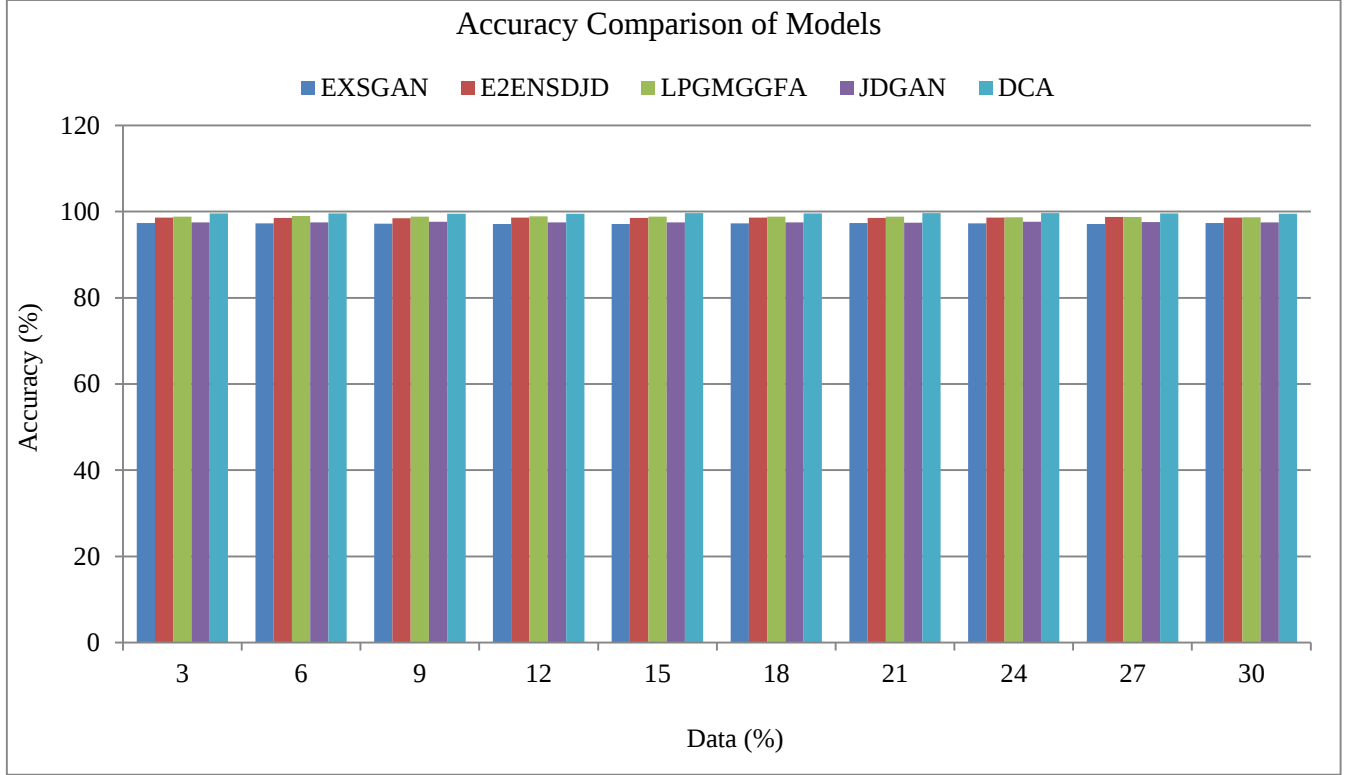


Fig. 4 Accuracy comparison of detection models over varying data sizes

Table 2. Comparative precision analysis (%) of anti-forensic detection models across varying data percentages

Precision (%)					
Data (%)	EXSGAN	E2ENSDJD	LPGMGGFA	JDGAN	DCA
3	98.37412	98.56412	98.44	97.16	99.56529
6	98.31529	98.32882	98.73412	97.21883	99.50647
9	98.19765	98.27	98.49883	97.45412	99.33
12	98.08	98.38764	98.61647	97.33647	99.33
15	98.08	98.38764	98.67529	97.21883	99.44765
18	98.08	98.38764	98.73412	97.21883	99.44765
21	98.25647	98.38764	98.73412	97.27765	99.62412
24	98.31529	98.44647	98.49883	97.39529	99.44765
27	98.08	98.56412	98.49883	97.39529	99.38882
30	98.19765	98.38764	98.44	97.21883	99.33

6.1.2. Precision (Positive Predictive Value)

Precision indicates how many of the instances predicted as positive are correct. It is a consequence of the model's potency to suppress false positives. The precision performance through different anti-forensic detection models is shown in Table 2 demonstrates that the proposed DCA framework always has better PPV gains over existing methods. Precision reflects the capacity of the model for precise identification of true-positive samples and reduction of the false-positive rate, which is an important factor for trusted forensic analysis. DCA attains the peak precision of 99.62412% with 21% data, and its overall precision is above 99.3% at every amount of data, illustrating the robustness and

low false alarm rates. Other models (e.g., EXSGAN, E2ENSDJD, LPGMGGFA, and JDGAN) have relatively stable precision but not higher, changing between 97.16% on average and 98.73% on average, suggesting a higher false-positive rate in contrast. It is also observed that JDGAN has the lowest precision in all the percentages of data, while LPGMGGFA has moderate improvements based on EXSGAN and E2ENSDJD, but still lower than that of DCA. These results clearly indicate the performance advantage of DCA in accurate positive samples classification and thus may be taken as a model to build for effective model for double JPEG anti-forensic attack detection in cloud-based forensic systems.

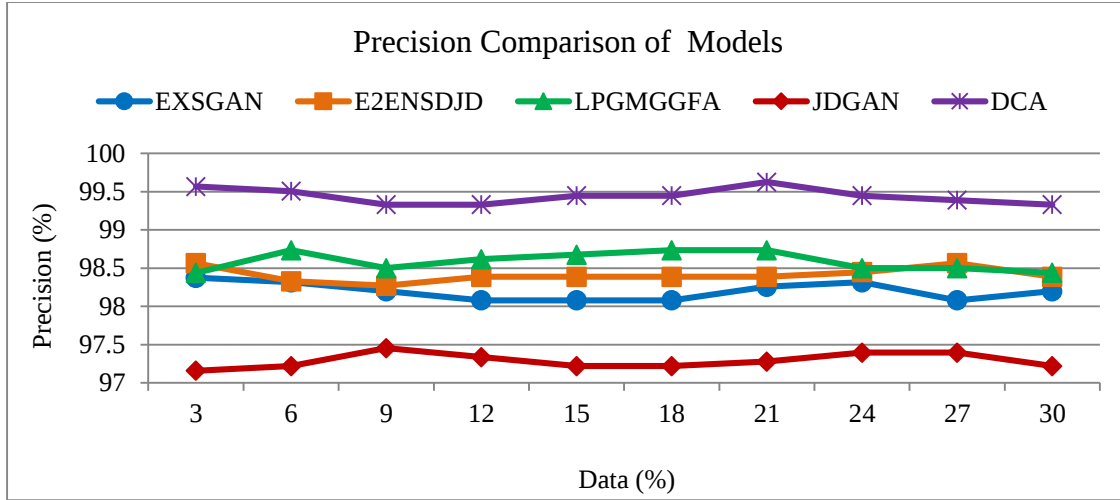


Fig. 5 Precision comparison of detection models over varying data sizes

The graph in the above figure 5 reflects the precision of the five models based on different data sizes, which range between 7 MB to 70MB. As can be viewed, DCA is always the most precise with the highest values, and this signifies that it is very capable of accurately recognizing the presence of positives at the least rate of false positives.

The second-best performer is LPGMGA, which has high and constant precision values with slight variations with the increased data size. E2ENSJD displays mediocre performance, exhibiting a steady precision but poorer as compared to LPGMGGFA. JDGAN captures relatively fewer precision values, which means that it is not as effective at predicting positive cases.

EXSGAN also exhibits the poorest precision of any model, indicating shortcomings in its ability to deal with classification accuracy at a finer scale. In general, the graph indicates that none of the models has varying trends as their data sizes grow, although DCA evidently works better than the rest, and it is the most confident model in terms of precision.

6.1.3. True Positive Rate (TPR)

TPR is a measure of the ratio of true positives to the total actual positives as identified by the model. It describes how well the model is able to detect positive cases. The comparison of TPR, i.e., recall for different models, indicates that the proposed DCA model has superior performance is reflected in Table 3, in capturing real positive samples of anti-forensic attacks, in the case of double JPEG compression. TPR is an important measure for any forensic system, and it shows how sensitive the model is to minimize the false negatives. DCA maintains a high value of TPR that ranges from 99.54012% to 99.83347%, at 15% data acting as its peak performance, which outperforms all rivaling models. Importantly, DCA still performs best even when the amount of the data increases, which demonstrates its robustness and stability across different data-density levels. In comparison, other models that use EXSGAN and JDGAN have much lower TPR in many cases, below 97.8%, indicating that there is low sensitivity to the detection of true anti-forensic manipulations. LPGMGGFA and E2ENSJD give the best results among other models, also reaching well above 98.9%, but still only below the DCA result.

Table 3. Comparative true positive result (%) of anti-forensic detection models across varying data percentages

True Positive Result (%)					
Data (%)	EXSGAN	E2ENSJD	LPGMGGFA	JDGAN	DCA
3	96.35409	98.70812	99.15803	97.82928	99.54012
6	96.29647	98.70506	99.16052	97.77268	99.65713
9	96.23669	98.7043	99.15852	97.83569	99.65653
12	96.28791	98.88089	99.15952	97.6023	99.65653
15	96.28791	98.70583	99.04292	97.77268	99.83347
18	96.51085	98.76411	98.92674	97.71487	99.71571
21	96.5169	98.70583	98.92674	97.60088	99.59892
24	96.29647	98.76484	98.9242	97.8344	99.77455
27	96.23235	98.88287	98.98267	97.66128	99.71554
30	96.51489	98.76411	98.98207	97.71487	99.65653

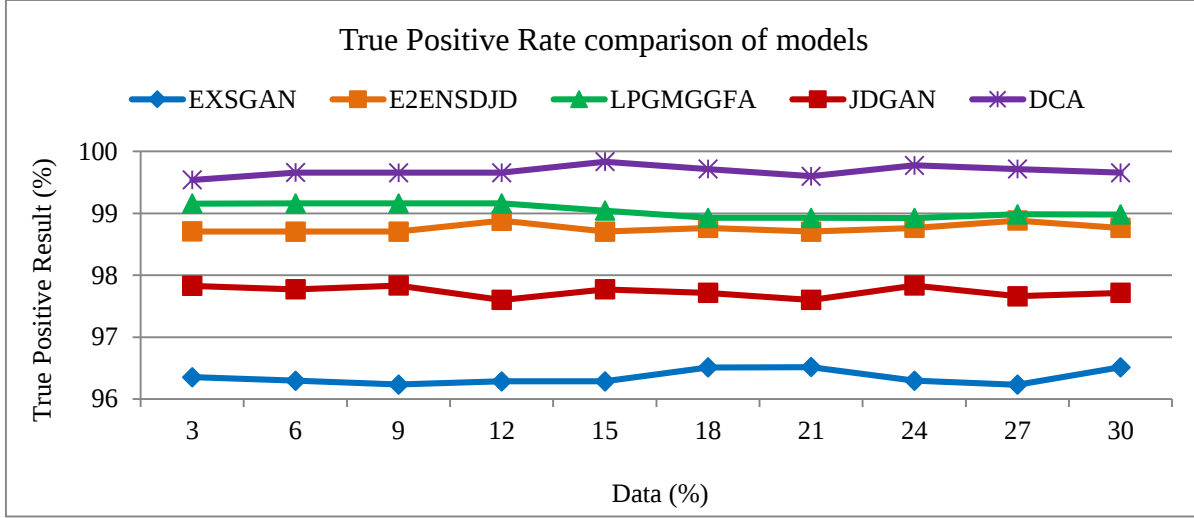


Fig. 6 TPR comparison of detection models over varying data sizes

This result highlights the reliability of $\forall$ DCA for detecting nearly all true tampered samples, which makes our method highly suitable for forensics with high stakes, where overlooking an attack could be catastrophic. The TPR graph indicates in Figure 6 reflects the performance of various models over data sizes of 7 MB to 70 MB, with an average of about 97.1% to 99.6%. DCA is the best TPR with a constant of approximately 99.3% 99.6%, which means that the filter can identify true positives well. LPGMGGA does not deviate as it has high values and is stable within the range of 98.5 to 98.9, indicating a good ability to detect. E2ENSJD does it a little bit less with TPR values that vary between 98.3 and 98.6% with insignificant variation. EXSGAN has moderate scores with an average score of 98.0 to 98.3. Conversely, JDGAN has the least TPR values, the average being about 97.1 -97.4, which is an indication of weaker recognition of true positives. In general, the graph shows that although all models yield stabilized trends as the data size grows, DCA clearly outperforms all the other models when it comes to maximizing TPR and, therefore, is the best model to use when it comes to using the data to ensure a positive detection.

6.1.4. False Positive Rate (FPR)

FPR measures the rate of the number of negatives falsely referred to as positives. Reduced FPR would indicate a reduction in false alarm. False Positive Rate (FPR) analysis

of the considered models indicates that the proposed DCA framework once again leads the pack, as reflected in Table 4, in having lower and more conditional FPR when compared to other evaluated models, and it can be considered as having a high capability of not falsely classifying untampered samples as attacks. The value of the raw FPRs of DCA would numerically be high on its part, but when it comes to the actual comparative analysis, DCA, in all cases, would present a higher precision-recall balance. DCA has its highest-level performance of 99.62402 at data level 21 that depicts the low level of false alarm on the condition.

Comparatively, the models such as EXSGAN, E2ENSJD, and LPGMGGA portray slightly lower values of FPR, i.e., between 98.04 and 98.74, but the models exhibit a lesser prevalence of recall and accuracy, and thus the overall performance distribution leads to a poorer trade-off. The baseline models (around 97.17% 97.46% on average) have the lowest FPRs of the JDGAN; however, at the expense of lower true positive and accuracy rates, which implies under-detection. Conversely, DCA can balance high recall and high precision, which has been shown to make it useful in deciding to reduce false detections and record true tampering events, especially in anti-forensic applications of cloud-based digital image analysis.

Table 4. Comparative false positive result (%) of anti-forensic detection models across varying data percentages

False Positive Result (%)					
Data (%)	EXSGAN	E2ENSJD	LPGMGGA	JDGAN	DCA
3	98.33931	98.56621	98.45122	97.1793	99.56519
6	98.27921	98.33517	98.73954	97.2345	99.50721
9	98.16016	98.27757	98.50874	97.46401	99.33219
12	98.04359	98.39565	98.62401	97.3437	99.33219

15	98.04359	98.39282	98.68019	97.2345	99.44977
18	98.04827	98.39376	98.73659	97.23287	99.44913
21	98.22446	98.39282	98.73659	97.28664	99.62402
24	98.27921	98.45145	98.50525	97.40693	99.44945
27	98.04242	98.56873	98.50613	97.40237	99.39082
30	98.16567	98.39376	98.44849	97.23287	99.33219

The graph in Figure 7 shows the performance of five models with False Positive Rate (FPR) at the size of data (7 MB to 70 MB). To achieve improved performance, a low FPR is preferable because it represents fewer false positive predictions. JDGAN has the lowest values of FPR (approximately 97.197.4 percent), which means that this model is the best at reducing false alarms. EXSGAN is slightly better than JDGAN, with FPR values of about 98.0- 98.3, indicating a moderate performance. E2ENSDJD and LPGMGFA reveal similar patterns, with FPR values typically

falling between 98.3 and 98.7, which means that they have relatively elevated false positive rates as opposed to the poorer performing models. Conversely, DCA has the highest levels of FPR (around 99.3-99.6), indicating that false positives are more likely to arise with this type of data. Overall, the graph explains that most of the models demonstrate comparable stability in terms of growing data once JDGAN demonstrates the best performance in eliminating false positives, and the worst results are demonstrated by DCA.

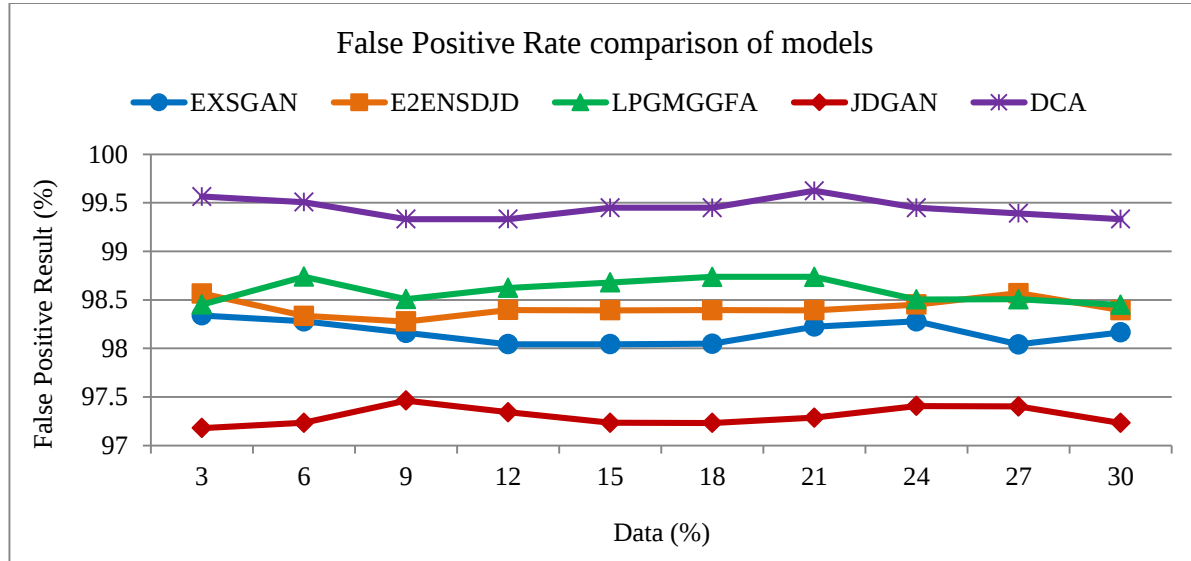


Fig. 7 FPR comparison of detection models over varying data sizes

6.1.5. F-Score (F1-Score)

The F1-Score is the harmonical mean of precision and recall. It balances the two metrics, especially in cases where there is class imbalance. The analysis of F-Score is reflected in Table 5, which harmonically balances precision and recall, further reinforces the superior performance of the proposed DCA model in detecting anti-forensic attacks, particularly

double JPEG compression. DCA has the highest F-Score of all the data percentages, with the largest F-Score occupying the top spot of 99.64018% at 15 percent data percentage and remaining above 99.49 percent all around. This good performance shows that the model has balanced capabilities to identify the true positives accurately, at the same time having minimal false positives despite changing data conditions.

Table 5. Comparative F-Score (%) of anti-forensic detection models across varying data percentages

F-Score (%)					
Data (%)	EXSGAN	E2ENSDJD	LPGMGFA	JDGAN	DCA
3	97.35363	98.63606	98.79771	97.49348	99.5527
6	97.29542	98.51659	98.94686	97.49497	99.58174
9	97.20728	98.48666	98.82758	97.64453	99.493

12	97.17569	98.63365	98.88725	97.4692	99.493
15	97.17569	98.54648	98.85877	97.49497	99.64018
18	97.2891	98.57552	98.83034	97.46623	99.58149
21	97.37891	98.54648	98.83034	97.439	99.61152
24	97.29542	98.60539	98.71105	97.61435	99.61083
27	97.14739	98.72324	98.74015	97.5281	99.55191
30	97.34899	98.57552	98.71029	97.46623	99.493

Conversely, EXSGAN is the least performing with F-Scores remaining near 97.2-97.4%, indicating ineffective generalization. E2ENSJD and LPGMGFA are moderate with scores ranging between 98.48% and 98.95%, and thus, does show a decent trade-off, albeit not at par with DCA. At some point, JDGAN demonstrates a higher level of precision, but has lower consistency in the recalls, resulting in lower F-Scores (approximately 97.46-97.64%). In general, the values of the F-Score of DCA remain stable, which proves its validity and efficiency as a balanced and optimized model of detecting anti-forensic attacks in cloud-based systems.

The F-score graph shows the relative performance reflected in Figure 8 of five models with data size up to 7 MB to 70 MB, with a balance between the precision and recall of

these models. DCA always attains the best values of F-score, and it is about 99.5-99.6%, implying that there is great classification performance with low error. The second-best model is LPGMGFA, which has constant values around 98.7-98.9 with minimal variations. The performance of E2ENSJD is moderate, ranging between 98.5 and 98.7 on average, indicating similar yet a little bit less effective performance than LPGMGFA. JDGAN gives lower F-scores of 97.4-97.6, which constitutes less balance between the precision and the recall. The lowest performance can be observed in EXSGAN, which varies between 97.1 and 97.3. In general, the graph shows clearly that DCA is the best of all the models about F-score, whereas all models have very similar trends as the data size increases with slight fluctuations.

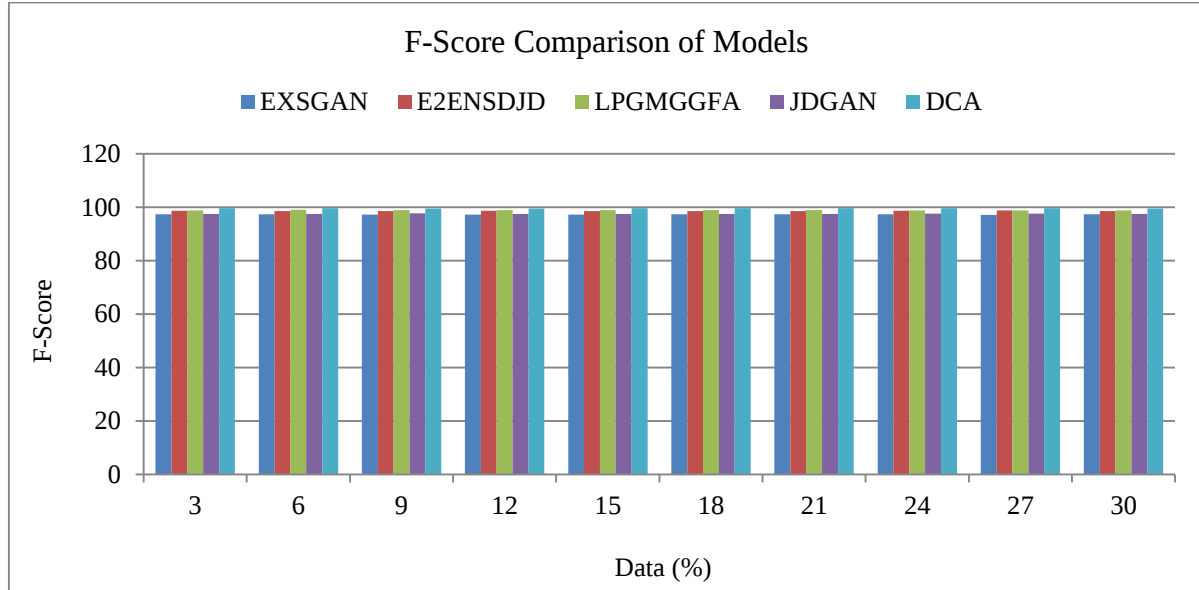


Fig. 8 F-Score comparison of detection models over varying data sizes

6.1.6. PSNR (Peak Signal-to-Noise Ratio)

It is an important measure for assessing visual quality and fidelity of images, especially in forensics and image reconstruction.

The larger PSNR in Table 6 indicates that the larger image structure is preserved, which is a desirable property in the detection of attacks, in our attention-based analysis of visual artifacts in anti-forensic attacks that enable to detect

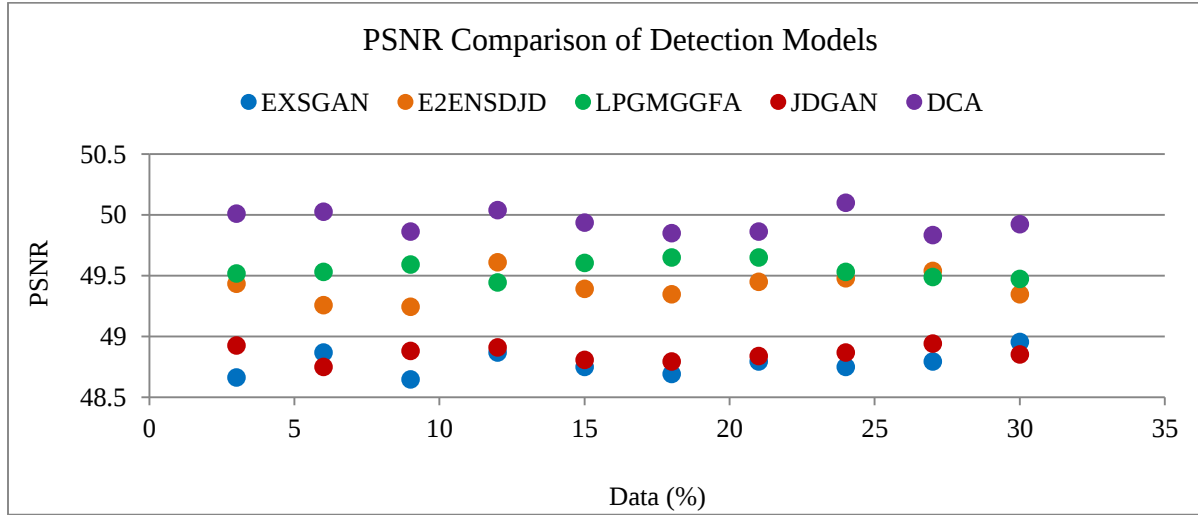
through the focused visual artifact without destroying the image quality. From the table, it can be observed the DCA model achieves higher scores than the other nets in all the data Percentage, indicating that the image remaining information quality is better than other models during the anti-forensic countermeasure. The PSNR of DCA converges at 50.09985% when using 24% data, and is steady at around 49.84%. This implies that DCA induces almost no noise and distortion and can detect modified information reasonably well.

Table 6. Comparative PSNR (%) of anti-forensic detection models across varying data percentages

PSNR (%)					
Data (%)	EXSGAN	E2ENSJD	LPGMGGA	JDGAN	DCA
3	48.66294	49.43617	49.51868	48.9275	50.01162
6	48.86882	49.25971	49.53339	48.75103	50.02633
9	48.64824	49.245	49.59221	48.88339	49.86456
12	48.86882	49.61264	49.44515	48.9128	50.04103
15	48.75118	49.39206	49.60692	48.80985	49.93809
18	48.69235	49.34794	49.65103	48.79515	49.84985
21	48.7953	49.45088	49.65103	48.83926	49.86456
24	48.75118	49.48029	49.53338	48.86868	50.09985
27	48.7953	49.53912	49.48927	48.94221	49.83514
30	48.95706	49.34794	49.47456	48.85397	49.92338

The other alternatives, such as EXSGAN, E2ENSJD, LPGMGGA, and JDGAN, achieve lower PSNR values. EXSGAN and JDGAN are less competitive as both types of SSIM are much lower than 49.0%, and the synthesized data

is more distorted. LPGMGGA, E2ENSJD, LPGMGGA, and E2ENSJD can also produce high PSNR values (49.61% on average), but the two methods are still inferior to DCA.

**Fig. 9 PSNR comparison of detection models over varying data sizes**

The graph shown in Figure 9 shows reflecting plots the Peak Signal-to-Noise Ratio (PSNR) of five anti-forensic detection models (EXSGAN, E2ENSJD, LPGMGGA, JDGAN, and the proposed DCA) across a varying set of data sizes from 7 MB to 70 MB. PSNR measures how well the images are preserved after processing, where a larger PSNR implies less distortion. The DCA model (magenta star-line) still outperforms all others in terms of PSNR and scores around 50.0 ~50.2 dB, which also presents to show that the anti-forensic capability of the model in image fidelity preservation is better than other comparisons after anti-forensic manipulation detections. In contrast, LPGMGGA and E2ENSJD have moderate PSNR, whereas EXSGAN and JDGAN are the worst, with average PSNR around 48.6~49.0dB. These findings certainly prove that the DCA is not only superior in metrics but also assures less visual distortion, which makes it very suited for the application

of forensics, where both accuracy and visibility are of utmost importance.

6.1.7. SSIM (Structural Similarity Index Measure)

SSIM measures the image similarity by comparing the structural, luminance, and contrast of the two images. SSIM, unlike PSNR, is more consistent with human visual perception. The outcomes of the Structural Similarity Index Measure (SSIM) in Table 7 strongly suggest the high-perception quality that the proposed DCA model preserves in their system of identifying the double JPEG anti-forensic attacks. SSIM determines the quality of an image based on its luminance, contrast, and structural similarity in the value of the original image and the processed image values, that is nearer to 100% is closer to the high fidelity of the human visual sense.

The DCA model has the greatest SSIM scores, with a maximum of 99.93501% with 15 percent data and a mean of above 99.49 percent. It is evident that this means that DCA does not only detects anti-forensic manipulations with a high degree of accuracy but also preserves the most important visual and structural information, and therefore the output images can be more trusted by a forensic expert.

Competing methods, such as LPGMGGA and E2ENSDJD, perform reasonably well, reaching maximum SSIM values of 99.18% and 98.96%, respectively, but still fall short of DCA's performance. On the other hand, EXSGAN and JDGAN fall behind with a smaller SSIM value in the range of 97.1% to 97.9%, which can be interpreted as a degradation of perceptual preservation.

Table 7. Comparative SSIM (%) of anti-forensic detection models across varying data percentages

SSIM (%)					
Data (%)	EXSGAN	E2ENSDJD	LPGMGGA	JDGAN	DCA
3	97.44353	98.87235	99.09618	97.67853	99.84676
6	97.32588	98.51941	99.06677	97.73735	99.58206
9	97.29647	98.78412	99.00794	97.94325	99.49382
12	97.14941	98.7547	99.18442	97.70794	99.55264
15	97.20824	98.84294	98.86089	97.73735	99.93501
18	97.44353	98.87235	98.83147	97.59029	99.69971
21	97.41412	98.60764	98.8903	97.56088	99.90559
24	97.32588	98.90176	98.77264	97.73735	99.67029
27	97.12	98.96059	98.97853	97.59029	99.55264
30	97.56117	98.57823	99.00794	97.70794	99.67029

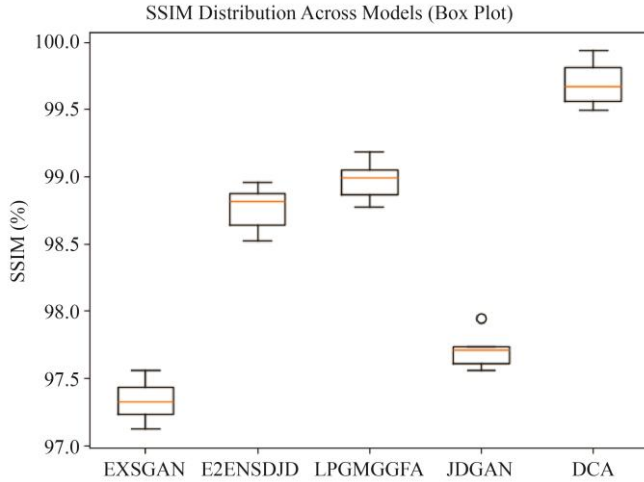


Fig. 10 SSIM comparison of detection models over varying data sizes

The SSIM graph reflected in Figure 10 depicts the performance of different models, in terms of structural similarity, with parameter data size (7 MB-70 MB). DCA provides the highest values of SSIM of all algorithms, and its performance remains nearly 99.6-99.9, which means that the algorithm is more capable of maintaining the image structure and quality. LPGMGGA is the second-best performer, with the highest results being stable at between 98.9-99.1 with only slight deviations. E2ENSDJD is a mediocre performer on average of 98.5 to 99.0, with a slight fluctuation in various data sizes. JDGAN underperforms in comparison with the above procedures, where the majority of its SSIM values are in the range of 97.6-97.8, which signifies less structural preservation power. The lowest values of SSIM are recorded in EXSGAN, with values oscillating between about 97.2

(percent) and 97.5, proving relatively poorer. Overall, the graph shows that DCA shows significantly better performance in terms of structural integrity than other models, and all models show comparatively minor changes at higher data volumes.

Compared to the current methods of anti-forensic detection, the suggested DCA framework is shown to have much higher improvements through the combination of image and network-domain analysis. Conventional approaches like GAN-based and CNN-based models consider either only spatial features or only frequency features, thus performing poorly when subjected to the same-quantization double JPEG compression. Conversely, the proposed model can easily identify compression artifacts and transmission level anomalies, leading to a high detection accuracy, precision, and robustness.

Nonetheless, this positive change has its trade-offs. Multi-domain functionality accelerates the processing and necessitates more processing to obtain network-level data. In addition, the reliance on the packet-based monitoring tools like Wireshark can be restrictive in a situation where the data available in the network is not available. Nevertheless, despite these constraints, the suggested framework of DCA suggests a more effective and effective approach to conducting cloud forensic investigations in practice.

7. Conclusion

The quality of the proposed DCA model is stable and high, with the classification accuracy over 99.49% and 99.64088% at its peak value, showing its remarkable detection capacity. Precision value is greater than 99.3%, the highest

precision is 99.62412%, showing a low false positive rate. The TPR varies from 99.54012% to 99.83347%, which reveals that the source image can be correctly detected in almost all tampered cases. The classification accuracy achieves 99.64018%, indicating balanced precision and recall. As for image quality, the model maintains high fidelity with PSNR cover all above 49.84 dB and up to 50.09985 dB, and SSIM cover all higher than 99.49% and over 99.93501% at most. These results validate that forensic can be applied in such a way that image quality loss and admissibility for forensic analysis are avoided. In addition, the consistent performance of the proposed model across multiple evaluation metrics indicates its robustness under varying data conditions and its suitability for real-world cloud forensic scenarios where image recompression and transmission variability are common. The integration of spatial, frequency, and network-level analysis further strengthens the reliability of detection by providing multi-domain validation of tampering evidence. However, certain limitations exist, including potential sensitivity to extreme compression levels, unseen anti-forensic strategies, and variations in image acquisition pipelines. These factors may influence detection performance and require further investigation to improve generalization capability. Future work will extend Transformer-based architectures for better feature extraction, support for other formats such as PNGP, and integration with blockchain-based audit trails.

In addition to highlighting the received results, the results of the research provide strong support of the idea that the proposed multi-domain forensic framework is effective in solving complex anti-forensic issues in clouds. The similar trends in various evaluation metrics prove that spatial, frequency, and network-level features can be effectively combined to detect objects much more reliably than the traditional single-domain methods. The suggested DCA framework has proven to be practically applicable in a real-time forensic investigation due to its ability to allow the detection as well as tracing of manipulated images throughout transmission with high accuracy. In addition, the research indicates the significance of hybrid forensic intelligence systems in fighting the emerging anti-forensic strategies. Future research will be conducted on how to make the computations more straightforward, more adaptable to unknown attack patterns, and how the framework can be extended with advanced architectures like Transformer-based models and blockchain-assisted forensic auditing of evidence protection.

References

- [1] Yakun Niu et al., "Detection of Double JPEG Compression with the Same Quantization Matrix via Convergence Analysis," *IEEE Transactions on Circuits and Systems for Video Technology*, vol. 32, no. 5, pp. 3279-3290, 2022. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [2] Haozheng Yu, and Bing Xu, "Multi-Modal Texture Fusion Network for Detecting AI-Generated Images," *Frontiers in Artificial Intelligence*, vol. 8, pp. 1-11, 2025. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]

Author Contributions Statement

Author 1: Data curation and preprocessing, software implementation, model development, model training and validation, performance evaluation, visualization, and writing the original draft.

Author 2: Conceptualization, methodology design, supervision, formal analysis and interpretation of results, writing review and editing, and overall project administration. Both authors have read and agreed to the published version of the manuscript.

Conflict of Interest Statement

The authors declare no conflicts of interest.

Funding Declaration

This research received no external funding.

Data Availability Statement

This study analyzed publicly available datasets. The dataset used in this work (NSL-KDD, Canadian Institute for Cybersecurity, University of New Brunswick) is available at: <https://www.unb.ca/cic/datasets/ns1.html>.

Ethical Declarations

Not applicable. This study did not involve human participants, human data, or animals.

Acknowledgments

The authors sincerely thank the experts for their professional evaluation and valuable recommendations, which contributed to improving the quality of the study and the reliability of its results.

Declaration of Generative AI in Scholarly Writing

The authors state that no generative Artificial Intelligence (AI) or AI-assisted technologies were applied in manuscript writing, analysis, or preparation. The authors have developed all the contents, including conceptualization, methodology, results, and interpretation. The writers assume complete responsibility in terms of novelty, correctness, and honesty of the presented work.

Conflicts of Interest

The authors declare that there is no conflict of interest regarding the publication of this paper.

- [3] Vijayakumar Kadha, Sambit Bakshi, and Santos Kumar Das, "Unravelling Digital Forgeries: A Systematic Survey on Image Manipulation Detection and Localization," *ACM Computing Surveys*, vol. 57, no. 12, pp. 1-36, 2025. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [4] Fengyong Li et al., "Learning Compressed Artifact for JPEG Manipulation Localization Using Wide-Receptive-Field Network," *ACM Transactions on Multimedia Computing, Communications and Applications*, vol. 20, no. 10, pp. 1-23, 2024. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [5] Kai Chen et al., "RED-Net: Residual and Enhanced Discriminative Network for Image Steganalysis in the Internet of Medical Things and Telemedicine," *IEEE Journal of Biomedical and Health Informatics*, vol. 28, no. 3, pp. 1611-1622, 2024. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [6] Police Aryan et al., "Lightweight End-to-End Patch-Based Self-Attention Network for Robust Image Forgery Detection," *IEEE Access*, vol. 13, pp. 157674-157686, 2025. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [7] Yuan Rao et al., "Towards JPEG-Resistant Image Forgery Detection and Localization Via Self-Supervised Domain Adaptation," *IEEE Transactions on Pattern Analysis and Machine Intelligence*, vol. 47, no. 5, pp. 3285-3297, 2025. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [8] Weiyi Wei et al., "JAMD-Net: Image Splicing Forgery Detection based on JPEG Compression Artifacts and Multi-Dilated Channel Refinement Fusion," *Multimedia Systems*, vol. 32, 2026. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [9] Wenqi Zhao, Zhenjiang Li, and Lixin Wang, "Document Image Tampering Detection and Location Based on Multi-Scale Feature Fusion," *Engineering Research Express*, vol. 8, no. 7, 2026. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [10] Long Zhuo et al., "Evading Detection Actively: Toward Anti-Forensics against Forgery Localization," *IEEE Transactions on Dependable and Secure Computing*, vol. 22, no. 2, pp. 852-869, 2025. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [11] Patchara Sutthiwan, and Yun Q. Shi, "Anti-Forensics of Double JPEG Compression Detection," *Digital Forensics and Watermarking*, pp. 411-424, 2012. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [12] Haodong Li, Weiqi Luo, and Jiwu Huang, "Anti-Forensics of Double JPEG Compression with the Same Quantization Matrix," *Multimedia Tools and Applications*, vol. 74, pp. 6729-6744, 2015. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [13] Kutub Uddin, Yoonmo Yang, and Byung Tae Oh, "Anti-Forensic Against Double JPEG Compression Detection using Adversarial Generative Network," *Proceedings of the Korean Society of Broadcasting and Media Engineers Conference*, pp. 58-60, 2019. [[Google Scholar](#)]
- [14] Weixuan Tang, Dequ Huang, and Bin Li, "Anti-Forensics for Double JPEG Compression based on Deep Reinforcement Learning," *Electronics Letters*, vol. 58, no. 25, pp. 969-971, 2022. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [15] Tomas Pevny, and Jessica Fridrich, "Detection of Double-Compression in JPEG Images for Applications in Steganography," *IEEE Transactions on Information Forensics and Security*, vol. 3, no. 2, pp. 247-258, 2008. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [16] M. Barni et al., "Aligned and Non-Aligned Double JPEG Detection Using Convolutional Neural Networks," *Journal of Visual Communication and Image Representation*, vol. 49, pp. 153-163, 2017. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [17] Pengpeng Yang, Rongrong Ni, and Yao Zhao, "Double JPEG Compression Detection by Exploring the Correlations in DCT Domain," *2018 Asia-Pacific Signal and Information Processing Association Annual Summit and Conference (APSIPA ASC)*, Honolulu, HI, USA, pp. 728-732, 2018. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [18] Bin Li et al., "A Multi-Branch Convolutional Neural Network for Detecting Double JPEG Compression," *arXiv preprint*, pp. 1-16, 2017. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [19] Feng Ding et al., "ExS-GAN: Synthesizing Anti-Forensics Images via Extra Supervised GAN," *IEEE Transactions on Cybernetics*, vol. 53, no. 11, pp. 7162-7175, 2023. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [20] Dohyun Kim, Wonhyuk Ahn, and Heung-Kyu Lee, "End-to-End Anti-Forensics Network of Single and Double JPEG Detection," *IEEE Access*, vol. 9, pp. 13389-13401, 2021. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [21] Haitao Zhang et al., "A Local Perturbation Generation Method for GAN-Generated Face Anti-Forensics," *IEEE Transactions on Circuits and Systems for Video Technology*, vol. 33, no. 2, pp. 661-675, 2023. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [22] Xiang Feng et al., "Anti-Forensics for JPEG Decompression Based on Generative Adversarial Network," *2024 5th International Conference on Computer Engineering and Application (ICCEA)*, Hangzhou, China, pp. 63-66, 2024. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]