

Investigation of image forgery via adaptive CNN and enhanced error level analysis



Mohammad Alsaffar¹, Diao Mohammed Uliyan^{2,*}, Moham'd Al-Dlalah³, Bander Nasser Almousa¹

¹Department of Information and Computer Science, College of Computer Science and Engineering, University of Hail, Hail 81481, Saudi Arabia

²Department of Information Security, College of Computer Science and Engineering, University of Hail, Hail 81481, Saudi Arabia

³High Diploma of Education, Arts and Humanities College, Applied Science Private University, Amman, Jordan

ARTICLE INFO

Article history:

Received 3 April 2026

Received in revised form

10 August 2026

Accepted 17 August 2026

Keywords:

Image forgery detection

Enhanced error level analysis

Adaptive convolutional neural network

Image splicing

Spatial features

ABSTRACT

Detecting image forgery is a significant challenge in digital image forensics, especially given the increasing use of advanced image editing software. This study proposes a new method that integrates Enhanced Error Level Analysis (EELA) with an adaptive Convolutional Neural Network (CNN) model to improve the detection of forged regions in images. Unlike traditional methods such as standard ELA or Residual Pixel Analysis, our approach uses EELA to examine compression artifacts in both original and forged images. Enhanced ELA calculates and displays error-level variations within the image at the pixel level. These features are then used to extract spatial features through an adaptive CNN training model to discriminate tampered regions. The proposed method was thoroughly evaluated on image benchmarks, including CASIA v1.0, CASIA v2.0, the Columbia Image Splicing Detection Dataset, and DEFACIO/IMD. The technique shows strong potential for identifying splicing forgeries. The proposed framework demonstrated superior results and greater resilience to image forgery attacks.

© 2026 The Authors. Published by IASE. This is an open access article under the CC BY-NC-ND license (<https://creativecommons.org/licenses/by-nc-nd/4.0/>).

1. Introduction

The advancement of image editing software has made it remarkably simple to modify images in different ways that may deceive individuals. These alterations raise significant concerns about image forgery. This type of forgery violates image privacy and can affect public confidence in visual media. Image forgery influences individuals' perceptions and cognitive processes, reshapes social conventions, and may also have legal consequences. Detecting altered regions in images is critical for preserving digital information and maintaining trust in visual media (Uliyan et al., 2016a). Many conventional forensic techniques used to determine image forgery depend on two types of methods: Error Level Analysis (ELA) and Residual Pixel Analysis (RPA).

The ELA method tries to reveal inconsistencies in compression features across different regions of an

image, whereas the RPA method examines pixel-level differences. These approaches frequently encounter difficulties with sophisticated modifications and depend significantly on manual interpretation, which may lack consistency (Tram et al., 2025).

Current developments in deep learning methods, such as Convolutional Neural Networks (CNNs), present a feasible solution. CNNs can recognize intricate image patterns and extract spatial features that many previous methods may have missed. By combining ELA with an adaptive CNN, it is possible to utilize the strengths of both methods. For example, ELA highlights compression discrepancies, while CNNs provide accurate classification of altered regions in images (Dell'Olmo et al., 2025; Kaur, 2025; Sohail et al., 2025).

This paper proposes a new framework that utilizes an Enhanced Error Level Analysis (EELA) method to visualize detailed error maps of images. These error maps are then fed into a CNN model to distinguish between authentic and manipulated areas. Experimental results obtained using the CASIA v1.0 and v2.0 datasets show that our approach outperforms existing methods. Integrating EELA with deep learning provides a reliable method for identifying image tampering and reinforcing confidence in digital visual content. The ELA and

* Corresponding Author.

Email Address: d.uliyan@uoh.edu.sa (D. M. Uliyan)

<https://doi.org/10.21833/ijaas.2026.08.012>

Corresponding author's ORCID profile:

<https://orcid.org/0000-0002-0777-9729>

2313-626X/© 2026 The Authors. Published by IASE.

This is an open access article under the CC BY-NC-ND license

(<https://creativecommons.org/licenses/by-nc-nd/4.0/>)

RPA methods face similar difficulties: they often cannot detect minor variations within images and may fail to correctly identify the differences between real and fake regions in complex images. Furthermore, these methods rely on manual inspection of artifacts, which can vary across examiners and may be insufficient against sophisticated modification techniques.

In deep learning methodologies, specifically, a CNN model serves as an effective approach for automated tamper detection. However, the majority of existing approaches inadequately utilize error-level information, which is essential for precisely identifying altered regions.

This paper introduces an integrated system that combines Enhanced Error Level Analysis (EELA) with an adaptive Convolutional Neural Network (CNN) model. The primary objectives are: To provide error-level mapping that reveals small discrepancies between authentic and fabricated photographs. To train a CNN with these EELA maps of images to identify the altered regions. To evaluate the proposed method against conventional ELA and RPA techniques, demonstrating enhancements in detection accuracy, precision, and recall. To create a strong and adaptable system for improving digital picture forensics, which will increase the reliability of visual content.

This paper introduces a new, reliable method that combines traditional error-level analysis with deep learning techniques to achieve these goals. As a result, this method helps to identify tampering in digital photographs.

2. Literature review

Image manipulation is performed to deceive or persuade readers. Image modification is becoming increasingly common as powerful editing software becomes more accessible. Image modification can cause inaccuracies, misrepresentation, and legal issues. Current image alteration detection approaches use deep learning, ELA, and CNNs. Because they can extract complex image features, these methods are more accurate than older approaches. ELA and deep learning methods have been used to achieve this. ELA identifies image alterations by comparing compression levels and image quality. Deep learning trains a convolutional neural network to detect tampering patterns (Biswas et al., 2024).

While these technologies have demonstrated good results in identifying picture alteration, they do have certain drawbacks. One of the most difficult issues is detecting copy-move fraud, which occurs when a portion of one picture is placed on another. Another problem is detecting manipulation in modified photos that have been post-processed to remove signs of tampering. Furthermore, identifying certain sorts of manipulation, such as compression, splicing, and resampling, remains difficult for many current approaches (Verma et al., 2024). Deep learning approaches have significantly improved

picture tamper detection accuracy. These approaches may extract complicated information from photos, resulting in more accurate tamper detection. They have performed better than traditional handmade feature-based approaches. Furthermore, deep learning approaches may adapt and generalize effectively to various forgery attacks (Jain and Singh, 2023).

Choudhary et al. (2024) presented a deep learning approach for detecting image forgery by combining Convolutional Neural Networks (CNNs) with Error Level Analysis (ELA). Using a UNet-based encoder-decoder model, the method accurately localizes tampered regions, particularly in challenging cases like image splicing. Experiments on the CASIA 2.0 dataset show strong performance, achieving 0.917 accuracy and 0.935 precision, surpassing existing transfer learning methods. A graphical user interface was also developed to enable practical verification of image authenticity.

Vaishali and Neetu (2024) proposed a residual network-based deep CNN (DCNN) image forgery detection method. To fix the problem of exploding and fading gradients, they utilized the residual network with 101 layers and skip connections. Furthermore, the cyclical learning rate (CLR) hyperparameter is applied to optimize the ResNet-101 model.

Despite these advances, many methods still have a hard time when the tampering has been deliberately covered up. Simple things like adding a bit of blur, some noise, or recompressing the image can weaken the very traces that ELA and CNN-based detectors depend on. Splicing further complicates detection because forged regions can be blended smoothly into the background, making compression inconsistencies too subtle even for well-trained models to identify reliably.

Recent studies have attempted to address this limitation by combining traditional signal-processing techniques with deep learning methods (Uliyan et al., 2016b). One approach used Hellinger distance and noise estimation to identify suspicious regions before passing them to a neural network, improving robustness against JPEG compression and Gaussian noise attacks (Das et al., 2026). Another method, AMSEANet, adopted an edge-guided multi-scale framework that integrates spatial and frequency-domain features to improve resistance against noise and compression artifacts that commonly reduce detection accuracy (Yang et al., 2025). Despite these improvements, detecting manipulated images that have been intentionally refined after tampering remains a significant challenge. This limitation highlights the need for more sensitive error-level analysis combined with adaptive feature-learning techniques.

3. Proposed framework

The main goal of the suggested approach is to demonstrate the accuracy of detecting forged regions. The method starts with image collection,

followed by preprocessing the images and applying the EELA transformation. Finally, an adaptive CNN model is built and trained to classify these images based on EELA feature maps. Fig. 1 illustrates the proposed process of our methodology. The Enhanced ELA method contributes to digital image forensics by improving the sensitivity and accuracy

of forgery detection. Through the integration of multi-scale analysis, channel weighting, and noise normalization, it effectively maps altered pixels while minimizing false positives caused by inherent image noise. This process makes the Enhanced ELA method stronger and more reliable for both visual examination and automated forensic evaluation.

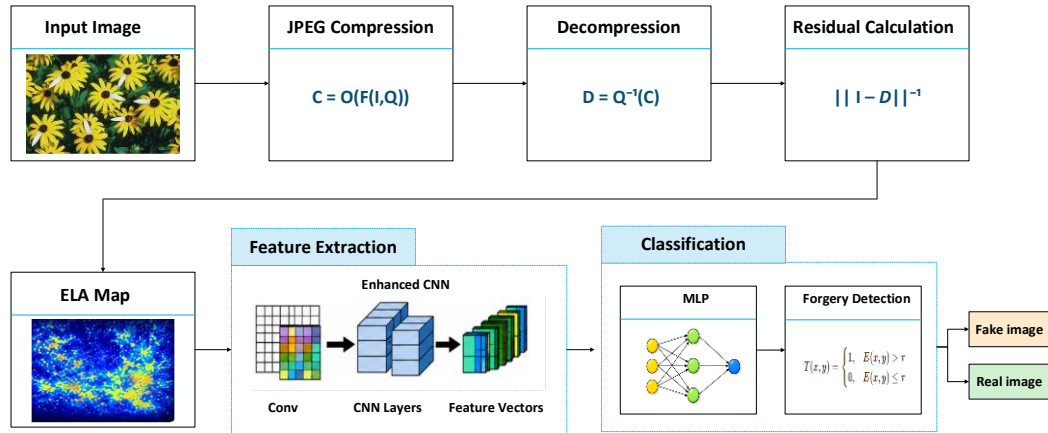


Fig. 1: Flowchart of the proposed method: EELA+ adaptive CNN

3.1. Data collection

The images are imported from the sourced image datasets CASIA v1, CASIAv2, Columbia Image Splicing Detection Dataset, and DEFACTO/IMD. It has a wide range of images with varying formats and features. The dataset is organized into two classes: real images and tampered images. Table 1 shows the information of images across these classes. The

dataset includes multiple image formats, such as JPG, TIF, and BMP, allowing the framework to be tested across different types of visual data. For ELA specifically, CASIA v2.0 and Columbia datasets are the most commonly used because they provide JPEG images at various quality levels, which ELA relies on. Fig. 2 and Fig. 3 show tampered images/original images.

Table 1: Information about benchmark dataset

Image dataset	Forgery type	Type of image	Size	Total number of original images	Total number of fake images	Notes for ELA
CASIA v1.0 (Dong et al., 2013)	Splicing	JPEG	~384×256	800 authentic images (all considered original)	921 tampered images (each forged from one or more originals)	Older, smaller dataset; can be used for benchmarking.
CASIA v2.0 (Dong et al., 2013)	Splicing, copy-move	JPG, TIF, BMP	~320x 240 to 800x600	7200 images	5123	Widely used; includes both authentic and tampered images; good for ELA because images are JPEG.
Columbia image splicing detection dataset (Hsu and Chang, 2006)	Splicing	TIFF/JPEG	~1280×960	180 authentic	180 spliced	Small but high-quality; perfect for testing ELA methods.
DEFACTO/IMD (Mahfoudi et al., 2019)	Splicing, copy-move, others	JPEG	~800×600 to 1600×1200	1,000+ images	~1,000	Contains metadata; some images include subtle manipulations suitable for compression artifact analysis.

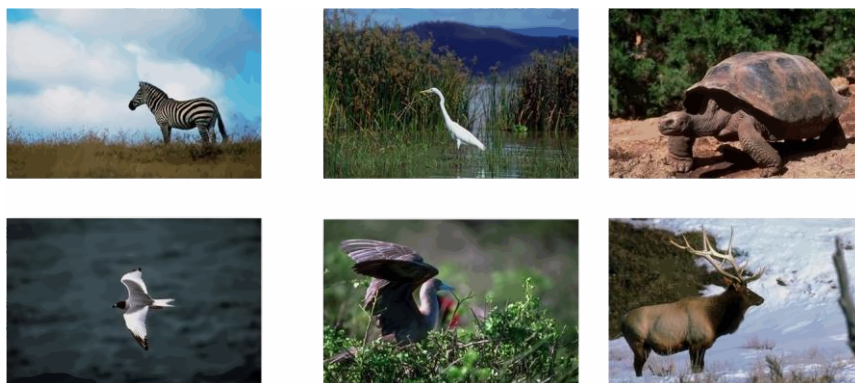


Fig. 2: Authentic images/CASIA v2 dataset

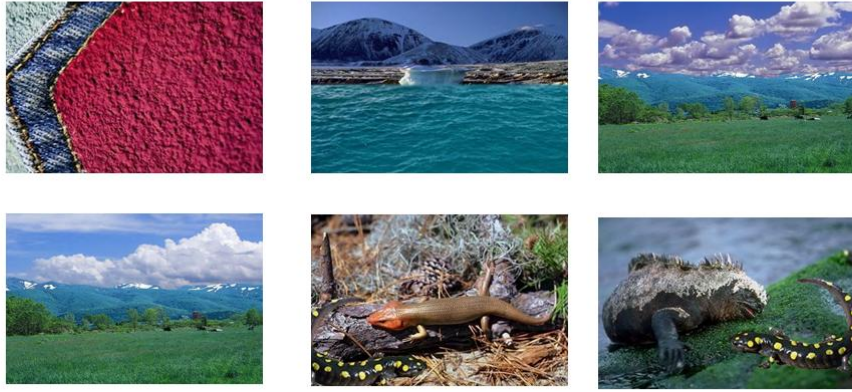


Fig. 3: Tampered images/CASIA v2 dataset

3.2. Image preprocessing

The preprocessing of the image is a crucial component in the preparation of images destined for the training of a deep learning model. It encompasses the following preprocessing procedures:

1. Size reduction: All photos are uniformly downsized to 128x128 pixels to provide consistent input dimensions.
2. Normalization: Adjust pixel intensities to a range of [0, 1] to perform sensible representation during training.
3. Image Expansion: It enhances the training dataset's size by techniques for example scaling and rotation.
4. Tuning: Adapting all images to RGB guarantees a uniform color space. Different types of images need to be normalized to RGB color for more straightforward and consistent processing in the following phases.

3.3. Enhanced error level analysis (EELA)

ELA is a method for examining images that may be used to show the changes between the original image and its compressed counterpart. This can help find digital image alterations. This is because various elements of the image compress differently depending on what they are and any changes made to the image when it is saved in a lossy compression format like JPEG. ELA essentially enables the visualization of discrepancies to identify areas of manipulation. ELA is a technique used to detect image tampering. It works on the principle that JPEG compression is lossy, meaning that each time you save a JPEG, small artifacts are introduced. ELA identifies differences in compression levels across an image. The basic steps of the ELA algorithm are:

1. Take the original image and resave it at a known JPEG quality (e.g., 90%).
2. Compute the difference between the original and resaved image.
3. The resulting difference image highlights regions with unusual compression, which may indicate manipulations, and detects two types of areas:

Uniform areas → likely untouched and Bright or highly contrasting areas → potentially edited.

Enhanced ELA improves traditional ELA by adding adaptive compression, noise-aware normalization, multi-scale analysis, channel-based weighting, and optionally deep learning classification. It highlights tampered regions in digital images by detecting compression inconsistencies. It consists of the following steps:

Step 1: Adaptive resaving

$$Q_r = Q_{\text{orig}} - \Delta Q \quad (1)$$

$$I_{\text{resave}} = \text{JPEG}_{\text{compress}}(I_{\text{orig}}, Q_r) \quad (2)$$

where, Q_{orig} = estimated original JPEG quality, ΔQ = small offset (5 – 10%) of Q_{orig} . Slightly changes the compression to expose inconsistencies. The Impact of the offset value is: Too small → residuals are weak; too large → may introduce artifacts that confuse the detector. The predicted value of Q_{orig} for a given source JPG image is about 75–95. It renders compression artifacts which an increased ΔQ means greater residuals in regions that have been modified. for example, $\Delta Q = 5-10$ for tiny artifacts and $\Delta Q = 10-15$ for enhanced perception. if $Q_{\text{orig}} = 85$ and the forger recompress it with $Q = 75$, then $\Delta Q = 10$.

Step 2: Core ELA residual map to represent local noise variance per pixel

$$E(x, y) = |I_{\text{orig}}(x, y) - I_{\text{resave}}(x, y)| \quad (3)$$

$$\text{Normalization is } E_{\text{norm}}(x, y) = \frac{E(x, y)}{\max(E(x, y))}$$

Step 3: Noise-aware normalization

$$E_{\text{norm}}(x, y) = \frac{E(x, y)}{N(x, y) + \epsilon} \quad (4)$$

where, $N(x, y)$ = local noise variance, ϵ = small constant to avoid division by zero. Better noise estimation improves differentiation between natural image texture and tampering artifacts.

Step 4: Multi-scale analysis

$$E_s = |I_{\text{orig}}^{(s)} - I_{\text{resave}}^{(s)}| \quad (5)$$

$$E_{\text{multi}}(x, y) = \sum_s w_s \cdot \text{Upsample}(E_s) \quad (6)$$

where, s = scale levels at which the image is analyzed, w_s = weight for each scale s to determine the influence of each scale on the final multi-scale map as shown in Fig. 4.

Step 5: Channel-based analysis can weight channels differently if tampering artifacts are more prominent in certain channels. It influences sensitivity to color-specific manipulation.

$$E_{\text{RGB}}(x, y) = \frac{1}{3}(E_R(x, y) + E_G(x, y) + E_B(x, y)) \quad (7)$$

Step 6: CNN learning integration

$$F = \phi(E_{\text{final}}; \theta_\phi) \quad (8)$$

$$P_{\text{tamper}}(x, y) = \sigma(\psi(F; \theta_\psi)) \quad (9)$$

where, ϕ = feature extraction layers, ψ = classifier layers, σ = sigmoid activation and $P_{\text{tamper}} \in [0, 1]$ = probability of tampering.

Step 7: Map P_{tamper} to a heatmap to highlight potential tampered regions. To convert probability maps into binary tamper masks, Typical Threshold = 0.5 is employed for equal sensitivity/specificity. Lower thresholds $\rightarrow$ more false positives; higher thresholds $\rightarrow$ risk of missing tampering.

After computing the ELA residual map $E(x, y)$ (or the final probability map $P_{\text{tamper}}(x, y)$ if using CNN), the values are continuous. Thresholding converts these continuous values into a binary map:

$$T(x, y) = \begin{cases} 1, & E(x, y) > \tau \\ 0, & E(x, y) \leq \tau \end{cases} \quad (10)$$

where, $1 \rightarrow$ pixel is suspected tampered, $0 \rightarrow$ pixel is likely authentic. Normal threshold values are set for Fixed: 0.05 to 0.1 (for normalized residuals, 0 to 1) and for Adaptive: Mean + $k \times \text{std}$ (where k is around 1–2). In CNN training to find EELA-based forgeries, the size of the batch depends on the memory of the GPU and the resolution of the image. Normal range: 16 to 32 for datasets and high-res images.

Fig. 4 shows the comparison of original and manipulated images with error-level analyses: (a) original image, (b) manipulated (fake) image, (c) ELA of the original image, and (d) EELA of the manipulated image, highlighting residual artifacts introduced by the spliced forgery.

3.4. Adaptive CNN model architecture

The deep learning model used for image tampering detection is a Convolutional Neural Network (CNN). The model is a consecutive sequence of layers where each layer has one input

vector and one output vector. Common Enhancements on CNN were employed as follows:

1. Residual Connections (ResNet) are used. Skip connections are added to ResNet to prevent vanishing gradients. Prevent vanishing gradients in deep networks by allowing the input to bypass (skip) convolutional layers. This makes training very deep networks stable and more effective. The input x is added to the output of a convolutional operation $F(x, W)$ forming a residual block.

$$y = F(x, W) + x \quad (11)$$

where, $F(x, W)$ is the convolutional operation, x is the input, and y is the output of the residual block.

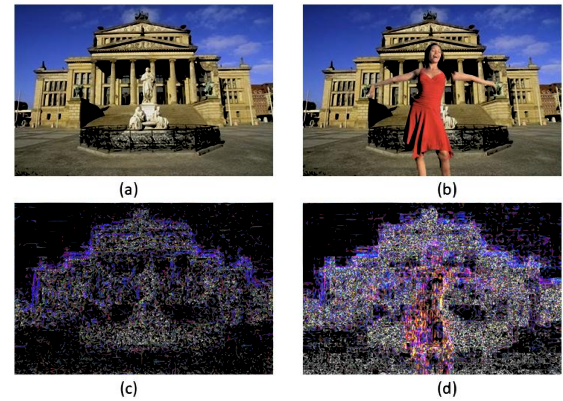


Fig. 4: (a) original image, (b) fake image, (c) ELA of image (a), (d) EELA of image (b)

2. Attention mechanisms (e.g., SE Block) focus on important features in spatial and channel dimensions. Channel Attention is defined as:

$$F_{\text{channel}}(X) = X \cdot \sigma(W_2 \delta(W_1 \text{GAP}(X))) \quad (12)$$

where, GAP = global average pooling, W_1, W_2 are learned weights, σ = sigmoid.

3. Multi-scale feature extraction captures fine and coarse details using parallel convolution filters (like Inception modules) defined as

$$F_{\text{multi}}(X) = [f_1(X), f_3(X), f_5(X), f_{\text{pool}}(X)] \quad (13)$$

F_{multi} concatenates feature maps of different kernel sizes.

4. Normalization layers (BatchNorm/LayerNorm) stabilize training and speed convergence:

$$\hat{x}_i = \frac{x_i - \mu_{\text{batch}}}{\sqrt{\sigma_{\text{batch}}^2 + \epsilon}} \quad (14)$$

Normalization adjusts the activations of a layer to have zero mean and unit variance, computed over a batch (BatchNorm) or layer (LayerNorm).

5. Dropout/spatial dropout reduces overfitting by randomly dropping units during training. For

Enhanced ELA, the enhanced CNN can take the ELA residual map as input:

$$P_{\text{tamper}} = f_{\theta}(E_{\text{final}}) \quad (15)$$

where, f_{θ} is the enhanced CNN with attention, residuals, and multi-scale features. More details on the implementation of experimental parameters for the adaptive CNN model are shown in Table 2. The proposed Adaptive CNN framework, as shown in Fig.

5, processes RGB, grayscale, and EELA-transformed images to detect image forgeries. Input images are first preprocessed using normalization and EELA transformation to generate residual error maps. EELA map images are fed into an adaptive CNN model that incorporates convolutional layers, residual blocks with SE attention, and pooling operations to extract discriminative feature vectors. Then features are classified using a fully connected MLP, which generates a real or fake prediction.

Table 2: Hyperparameters of adaptive CNN

Component	Configuration
Baseline models	SE block
Proposed model	Adaptive CNN, ResNet-SE, attention-based hybrids
Input	RGB, grayscale, EELA-transformed images
Loss	Binary cross-entropy (classification), pixel-wise (localization)
Optimizer	Adam/SGD with momentum
LR/batch/epochs	$1e^{-3}$ / 16-32 / 50-100
Regularization	Weight decay, dropout
Hardware	NVIDIA RTX 3090 GPU(s)
Framework	PyTorch/TensorFlow
Metrics	Accuracy, precision, recall, F1, AUC
Early stopping	Validation loss-based
Evaluation	Test set, cross-dataset
Visualization	Heatmaps, confusion matrix

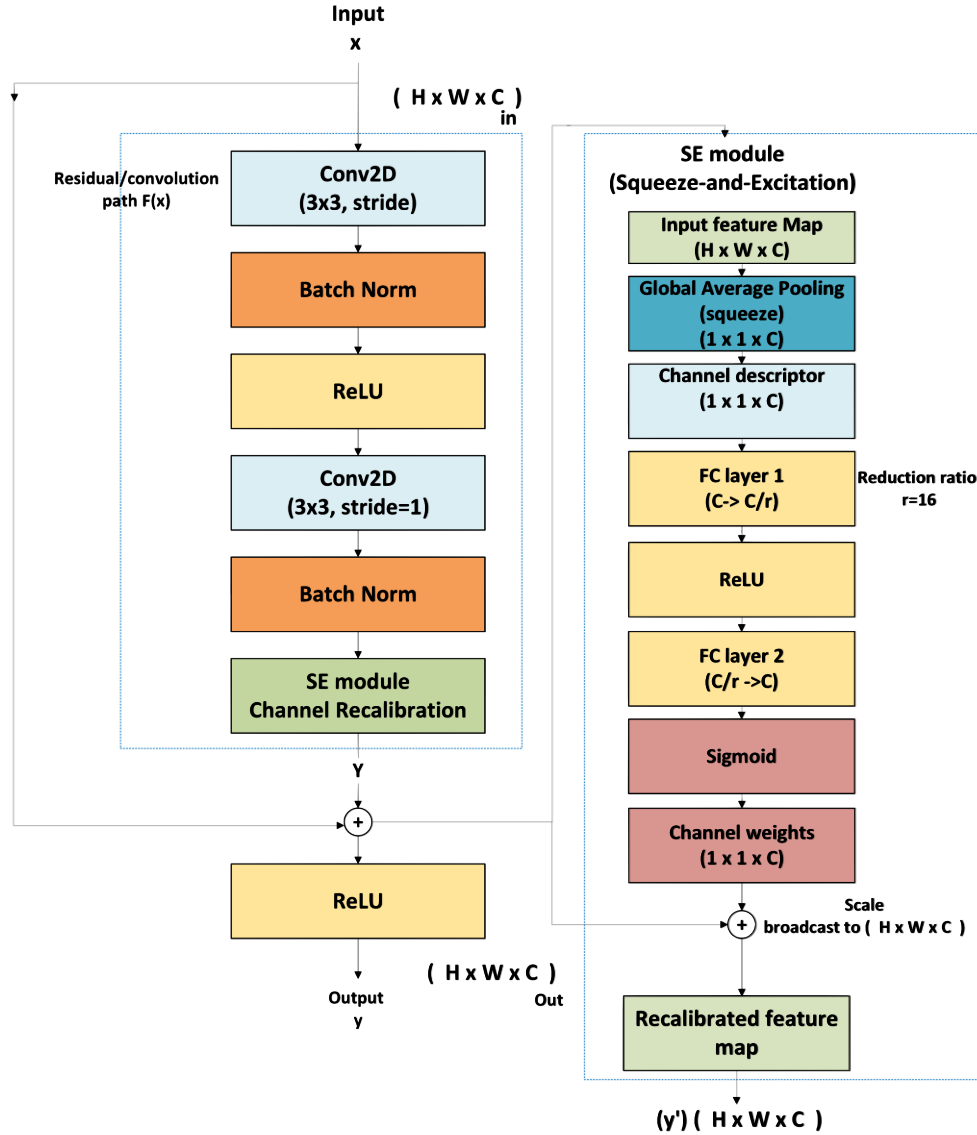


Fig. 5: Adaptive CNN framework

The concise pseudocode as described in Fig. 6 takes the residual path, SE module, and skip connection into the adaptive CNN model.

```

FUNCTION ResNet_SE_Block (InputFeatureMap X, INTEGER reduction_ratio = 16)
// 1. Residual / Convolutional Path F(x)
F = Conv2D (X, kernel_size = 3, stride = 1, padding = 1) // Conv layer 1
F = BatchNormalization (F) // Batch norm
F = ReLU (F) // Activation
F = Conv2D (F, kernel_size = 3, stride = 1, padding = 1) // Conv layer 2
F = BatchNormalization (F)
// 2. Squeeze-and-Excitation (SE) Module
S = GlobalAveragePooling (F) // a channel descriptor of size=1 X 1 X C
S = FullyConnected (S, units = C / reduction_ratio) // Dimensionality reduction
S = ReLU (S)
S = FullyConnected (S, units = C) // Restore channels
S = Sigmoid (S) // Channel-wise weights (0~1)
// Scale F by learned channel weights
F_se = F * S // Element-wise multiplication
// 3. Skip Connection
Y = F_se + X // Residual addition
Y = ReLU (Y) // Final activation
RETURN Y
END FUNCTION

```

Fig. 6: Pseudocode of adaptive CNN model

3.5. Training and validation

The model learns from the training dataset, and its performance is checked on a different dataset called the validation dataset. To put it another way, the training process has these steps:

- **Splitting the Data:** The data set is split into two parts: a training set and a validation set, with 80% of the data going to the training set and 20% going to the validation set.
- **Model Compilation:** The Adam optimizer is utilized with a learning rate of 0.001. The binary cross-entropy loss is what is employed here.

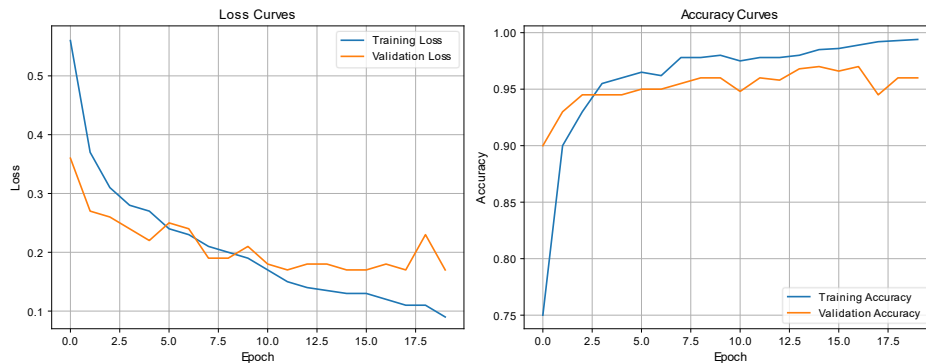


Fig. 7: Loss and accuracy curves of the proposed model

Accuracy is the ratio of successfully categorized photos to the total number of images.

$$Accuracy = \frac{(TP+TN)}{(TP+TN+FP+FN)} \quad (16)$$

Precision is defined as the ratio of true positive images to the total number of images that are classified as positives.

$$Precision = \frac{TP}{(TP+FP)} \quad (17)$$

Recall is defined as the ratio of true positive images to the total number of positive images.

- **Training:** 30 epochs were performed to keep it from overfitting.
- **Evaluation:** The model's performance is measured by accuracy, precision, recall, and the F1 score as shown in Fig. 7. It is also shown with a confusion matrix, which is a way to show how well the categorization works.

Fig. 7 illustrates the training and validation loss and accuracy across 20 epochs of the proposed model. Both training and validation losses diminish over time, signifying that the model is acquiring knowledge.

4. Results and discussion

The suggested method's efficiency is corroborated by the examination of the renowned CASIA V1.0,2.0 image-tampered database. There are 12,614 images in the database, and they are in four different formats: BMP, JPG, PNG, and TIF. Of these, 7,200 are original shots, and 5,123 have been changed. CASIA 2.0 shows several types of things, such as animals, buildings, articles, characters, plants, nature, landscapes, textures, and images of the inside of structures. The photos in this collection range in size and resolution from 800 × 600 pixels to 384 × 256 pixels. All the information about the CASIA 1.0, CASIA 2.0, Columbia Image Splicing Detection Dataset, and DEFACTO/IMD is described in Table 1. This section analyzes performance based on metrics such as accuracy, precision, recall, and the F1 score. To ascertain these figures, it is essential to comprehend the following terminology.

$$Recall = \frac{TP}{(TP+FN)} \quad (18)$$

F1 Score is defined as the harmonic mean of precision and recall.

$$F1 \text{ Score} = 2 \times \frac{Precision \times Recall}{(Precision+Recall)} \quad (19)$$

To provide a self-evaluation of the proposed image forgery detection framework, each image from the Columbia Image Splicing Detection Dataset was first preprocessed via Enhanced ELA to highlight regions with potential forgery attacks. These ELA-mapped images were then fed into an Adaptive

Convolutional Neural Network (CNN) designed to capture both local and global features, with attention mechanisms incorporated to focus on likely forged regions.

The Adaptive CNN was trained through the Adam optimizer with a learning rate of 1×10^{-3} , a batch

size of 16, and for 50 epochs. The efficacy of our system was evaluated using accuracy, precision, recall, and F1-score, thereby testing the proposed method through a thorough assessment of detection capabilities and resilience across many forgery types, as shown in Table 3.

Table 3: Performance comparison of different methods on the Columbia image splicing detection (CISD) dataset

Method	Accuracy (%)	Precision (%)	Recall (%)	F1-score (%)
Standard CNN (baseline)	88.5	87.2	88	87.1
Standard ELA + CNN	89.8	88.5	90.3	89.4
Adaptive CNN + Enhanced ELA	94.2	93.5	95	94.2

Fig. 8 shows the comparison between four image forgery detection methods: EELA+CNN model, ELA+VGG16, ELA+VGG19, and ELA+ResNet101. The suggested framework achieves an accuracy = 96.21%, greatly superior to the other models. VGG19 and VGG16 also perform effectively, exhibiting accuracy=90.32% and 88.92%, respectively. With 74.75% accuracy, ResNet101 is less accurate. This comparison shows that the EELA+CNN model improved its ability to identify altered regions in images.

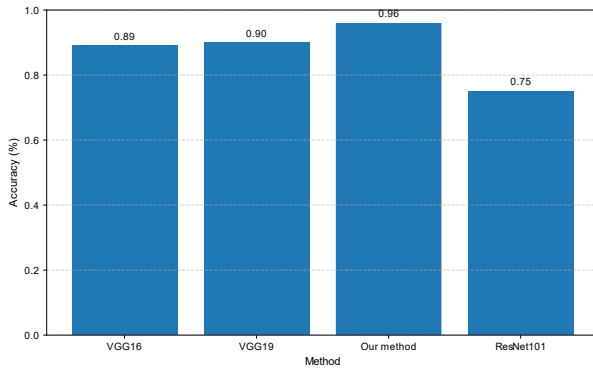


Fig. 8: Accuracy results

Fig. 9 shows precision, which measures the number of true positives compared to the total number of positive identifications made by the EELA+CNN model compared with other methods. The proposed model has a precision of 98.5%, which means it is quite reliable in finding altered photos. VGG16, on the other hand, has an accuracy of 92.8%, and VGG19 has an accuracy of 88%. Therefore, the precision score measures how well the suggested model reduces false positives.

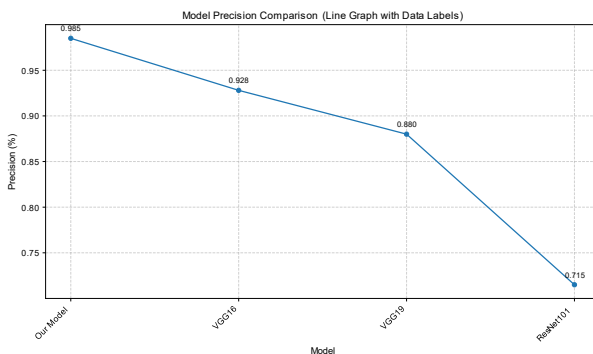


Fig. 9: Precision comparison

Fig. 10 displays the recall comparison bar chart and how well the models recognize manipulated dataset photos. The proposed model leads with 92.4% recall, detecting most manipulated photos. VGG19, VGG16, and ResNet101 have lower recall rates of 89.2%, 80.2%, and 66.8%. High recall in the suggested approach indicates robustness in spotting tampered areas.

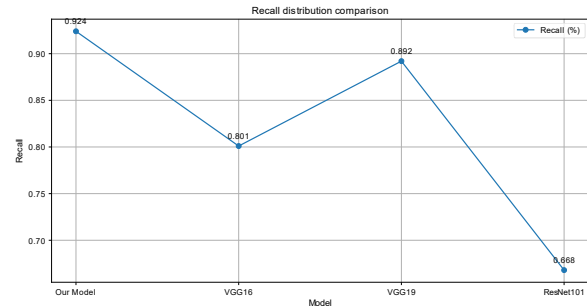


Fig. 10: Recall comparison

Fig. 11 shows the bar chart comparing the F1 score of four different models in detecting image tampering. The proposed model achieves the highest F1 score of 95.5%, significantly outperforming the others.

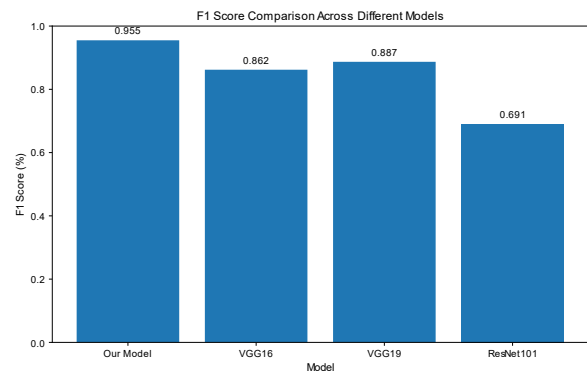


Fig. 11: F1 score comparison

VGG19 and VGG16 perform well with 88.70% and 86.2%. With a 69.1% F1 score, ResNet101 trails. The comparison suggests that the CNN model excels at identifying fabricated images. Table 4 presents a comparative analysis of four models, evaluated based on Accuracy, Precision, Recall, and F1 Score. The suggested model exhibits superior performance, achieving the highest values across all metrics: an accuracy of 96.21%, a precision of 98.58%, a recall of 92.36%, and an F1 score of 95.37%. While VGG16

and VGG19 demonstrate moderate performance, their recall and F1 scores are not particularly noteworthy. Conversely, ResNet101's metrics are suboptimal for this application. This comprehensive evaluation underscores the efficacy of the proposed methodology in identifying image manipulation.

By using CNN along with ELA, the confidence level of the real image is improved to 99.70%, and the fake image is improved to 99.89%, as shown in Fig. 12 and Fig. 13. Fig. 14 illustrates the confusion matrix, which shows how the model performs in terms of accuracy and precision.

Table 4: Comparison of performance metrics with other models

Model	Accuracy (%)	Precision (%)	Recall (%)	F1 score (%)
ELA+VGG-16 (Choudhary et al., 2024)	88.92	92.69	80.07	85.92
ELAResNet-101 (Vaishali and Neetu, 2024)	74.75	71.53	66.78	69.07
ELA+VGG-19 (Khalil et al., 2023)	90.32	88.16	89.04	88.60
Proposed model (EELA+CNN)	96.21	98.58	92.36	95.37

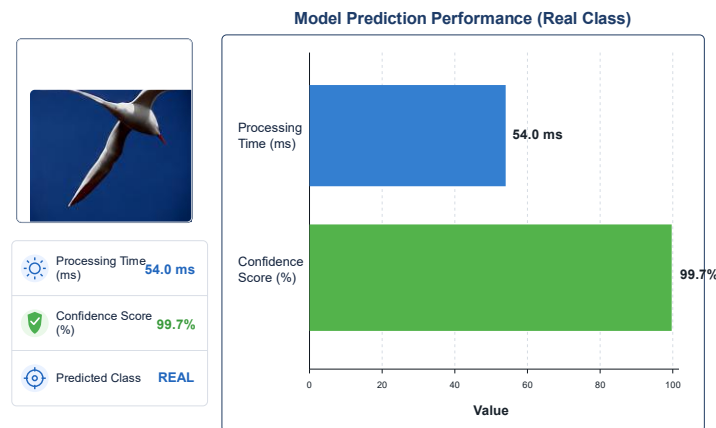


Fig. 12: Prediction and confidence of real image

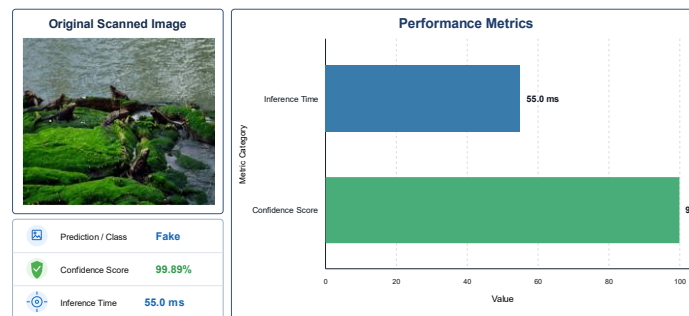


Fig. 13: Prediction and confidence of fake image

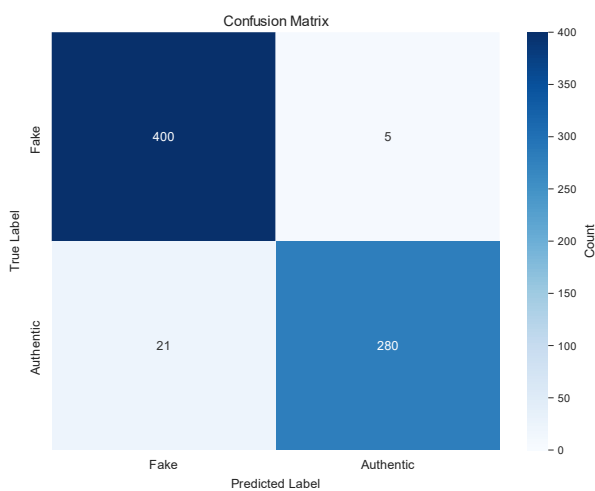


Fig. 14: Confusion matrix

Based on the confusion matrix provided, here are the professional comments about the results. High Identification Accuracy: The model demonstrates an overall Accuracy of 96.3% (680/706 samples). The

heavy concentration of values along the main diagonal (400 and 280) visually confirms that the model is highly effective at distinguishing between fake and authentic classes.

Superior Sensitivity to "Fake" Samples: High Recall (98.7%): The model successfully caught 400 out of 405 fake samples. Low Leakage: Only 5 fake samples were misclassified as authentic (False Positives). This is a critical strength if your goal is security or fraud detection, as it minimizes the risk of a "fake" item passing as "real."

Robust "Authentic" Classification: The model correctly identified 280 authentic samples. There were 21 instances where authentic samples were flagged as fake (False Negatives). While this is a small percentage (approx. 7%), it represents the primary area for potential improvement to reduce "false alarms" for legitimate users.

Balanced Performance: The F1-Score of 96.8% indicates an excellent balance between Precision and Recall. The model is not biased toward one class; it

maintains high reliability whether it is identifying a forgery or verifying an original.

Limitations of EELA + adaptive CNN method:
Subtle or Artifact-Free Manipulations: Some manipulations do not leave strong compression artifacts, making them difficult for ELA-based methods, as mentioned in Table 5.

Images of Low Resolution: In the case of small images with lower pixel resolution, CNN feature maps exhibit fewer spatial features. This may cause subtle forgeries to become undetectable in low-resolution residual maps. Pooling layers in convolutional neural networks may eliminate minor artifact patterns. For example, compared with a 512×512 image, slight variations within a 64×64 image may become imperceptible in CNN feature maps.

Intense JPEG Compression: When Q_{orig} is very low ($\leq 50-60$), an image exhibits significant artifacts. Residuals from modification are concealed by global compression distortion. Then, False positives may rise; the CNN may improperly classify genuine regions as counterfeit.

Table 5: Common forgery types and their detection challenges using ELA-based methods

Forgery type	Challenge
Resizing/scaling without recompression	Minimal compression artifact differences; ELA residuals may be weak
Copy-move within same JPEG quality	Local duplication produces weak residuals if compression is identical
High-quality splicing/seamless blending	Blending artifacts are smoothed; CNN may miss subtle features
Low-contrast tampering	Pixel changes are too small to generate distinguishable residuals
Double compression with same Q	$\Delta Q \approx 0 \rightarrow$ ELA residuals nearly vanish

5. Conclusion

This paper presented an image forgery detection framework that integrates Enhanced Error Level Analysis (EELA) with an adaptive Convolutional Neural Network (CNN). EELA was used to generate enhanced residual maps by incorporating adaptive compression, noise-aware normalization, multi-scale analysis, and channel-based analysis, while the adaptive CNN was designed to learn discriminative features from these maps for detecting manipulated regions. The experimental results demonstrated that the proposed EELA+CNN approach achieved an accuracy of 96.21%, precision of 98.58%, recall of 92.36%, and F1 score of 95.37% in the reported comparison. The method also achieved higher performance than the VGG-16, VGG-19, and ResNet-101 models evaluated in the study. These findings indicate that integrating enhanced error-level information with adaptive deep learning can improve the detection of image manipulation. Despite these results, the proposed approach has limitations. Manipulations that produce weak or no compression artifacts, low-resolution images, and images affected by strong JPEG compression may remain challenging to detect. Future work should therefore investigate the performance of the framework on a wider range of datasets and more

diverse manipulation types, particularly cases in which conventional error-level analysis provides weak forensic evidence.

List of abbreviations

AMSEANet	Attention-based multi-scale encoder attention network
AUC	Area under the curve
BMP	Bitmap
CASIA	Chinese Academy of Sciences Institute of Automation image dataset
CISD	Columbia image splicing detection dataset
CNN	Convolutional neural network
CLR	Cyclical learning rate
DCNN	Deep convolutional neural network
DEFACTO/IMD	Image manipulation detection dataset/image manipulation dataset
EELA	Enhanced error level analysis
ELA	Error level analysis
GAP	Global average pooling
GPU	Graphics processing unit
JPEG	Joint photographic experts group
MLP	Multilayer perceptron
PNG	Portable network graphics
ReLU	Rectified linear unit
ResNet	Residual network
RGB	Red, green, and blue
RPA	Residual pixel analysis
SE	Squeeze-and-excitation
SGD	Stochastic gradient descent
UNet	U-shaped neural network
VGG	Visual geometry group

Acknowledgment

This research has been supported by the University of Hail, Saudi Arabia.

Compliance with ethical standards

Conflict of interest

The author(s) declared no potential conflicts of interest with respect to the research, authorship, and/or publication of this article.

References

- Biswas S, Islam MA, Rout J, Mishra M, and Muduli D (2024). Passive image forgery detection through ensemble learning with ELA. In the 2024 International Conference on Computer, Electronics, Electrical Engineering & Their Applications (IC2E3), IEEE, Srinagar Garhwal, Uttarakhand, India: 1-6. <https://doi.org/10.1109/IC2E362166.2024.10827508>
- Choudhary RR, Paliwal S, and Meena G (2024). Image forgery detection system using VGG16 UNET model. Procedia Computer Science, 235: 735-744. <https://doi.org/10.1016/j.procs.2024.04.070>
- Das D, Naskar R, and Chakraborty RS (2026). Image splicing localization based on Hellinger distance and noise estimation through convolutional neural network and vision transformer. Digital Signal Processing, 168(Part B): 105559. <https://doi.org/10.1016/j.dsp.2025.105559>
- Dell'Olmo PV, Kuznetsov O, Frontoni E, Arnesano M, Napoli C, and Randieri C (2025). Dataset dependency in CNN-based copy-

- move forgery detection: A multi-dataset comparative analysis. Machine Learning and Knowledge Extraction, 7(2): 54. <https://doi.org/10.3390/make7020054>
- Dong J, Wang W, and Tan T (2013). CASIA image tampering detection evaluation database. In the 2013 IEEE China Summit and International Conference on Signal and Information Processing, IEEE, Beijing, China: 422-426. <https://doi.org/10.1109/ChinaSIP.2013.6625374>
- Hsu YF and Chang SF (2006). Detecting image splicing using geometry invariants and camera characteristics consistency. In the 2006 IEEE International Conference on Multimedia and Expo, IEEE, Toronto, ON, Canada: 549-552. <https://doi.org/10.1109/ICME.2006.262447> **PMid:16739096**
- Jain S and Singh D (2023). VGG16Unet: Effective passive model for image splicing forgery localization. Journal of Electronic Imaging, 32(4): 043015. <https://doi.org/10.1117/1.JEI.32.4.043015>
- Kaur N (2025). Hybrid image splicing detection: Integrating CLAHE, improved CNN, and SVM for digital image forensics. Expert Systems with Applications, 273: 126756. <https://doi.org/10.1016/j.eswa.2025.126756>
- Khalil AH, Ghalwash AZ, Elsayed HAG, Salama GI, and Ghalwash HA (2023). Enhancing digital image forgery detection using transfer learning. IEEE Access, 11: 91583-91594. <https://doi.org/10.1109/ACCESS.2023.3307357>
- Mahfoudi G, Tajini B, Retraint F, Morain-Nicolier F, Dugelay JL, and Pic M (2019). DEFACTO: Image and face manipulation dataset. In the 2019 27th European Signal Processing Conference (EUSIPCO), IEEE, A Coruna, Spain: 1-5. <https://doi.org/10.23919/EUSIPCO.2019.8903181>
- Sohail S, Sajjad SM, Zafar A, Iqbal Z, Muhammad Z, and Kazim M (2025). Deepfake image forensics for privacy protection and authenticity using deep learning. Information, 16(4): 270. <https://doi.org/10.3390/info16040270>
- Tram PN, Nguyen CT, Le CT, Le DT, and Huynh KT (2025). Combining ELA and frequency analysis for robust classification of authentic, spliced and diffusion image. In the 2025 RIVF International Conference on Computing and Communication Technologies (RIVF), IEEE, Ho Chi Minh City, Vietnam: 623-628. <https://doi.org/10.1109/RIVF68649.2025.11365152>
- Uliyan DM, Jalab HA, Abdul Wahab AW, and Sadeghi S (2016a). Image region duplication forgery detection based on angular radial partitioning and Harris key-points. Symmetry, 8(7): 62. <https://doi.org/10.3390/sym8070062>
- Uliyan DM, Jalab HA, Wahab AWA, Shivakumara P, and Sadeghi S (2016b). A novel forged blurred region detection system for image forensic applications. Expert Systems with Applications, 64: 1-10. <https://doi.org/10.1016/j.eswa.2016.07.026>
- Vaishali S and Neetu S (2024). Enhanced copy-move forgery detection using deep convolutional neural network (DCNN) employing the ResNet-101 transfer learning model. Multimedia Tools and Applications, 83: 10839-10863. <https://doi.org/10.1007/s11042-023-15724-z>
- Verma A, Pandey P, and Khari M (2024). ELA-Conv: Forgery detection in digital images based on ELA and CNN. In: Santosh K et al. (Eds.), Recent trends in image processing and pattern recognition: 213-226. Springer, Cham, Switzerland. https://doi.org/10.1007/978-3-031-53082-1_18
- Yang Y, He Y, Ma X, Lv W, Chen J, and Wang H (2025). AMSEANet: An edge-guided adaptive multi-scale network for image splicing detection and localization. Sensors, 25(20): 6494. <https://doi.org/10.3390/s25206494> **PMid:41157548 PMCID:PMC12568058**