



A NOVEL FRAMEWORK FOR ANTI-FORENSICS IN DIGITAL FORGERY LOCALIZATION

VANTHADPULA VISHNU¹, NARSAPURAM PAVAN², GUGULOTHU VENKAT³, MATTAPARTHI SAI GANESH⁴, MOHAMMED AQEEMUDDIN⁵

ASSISTANT PROFESSOR¹, UG SCHOLAR^{2,3,4&5}

DEPARTMENT OF CSE, NARSIMHA REDDY ENGINEERING COLLEGE (UGC- AUTONOMOUS) MAISAMMAGUDA (V), KOMPALLY,
SECUNDERABAD, TELANGANA-500100

ABSTRACT

Digital image forensics has become an important research area for identifying manipulated or forged images. Modern forensic techniques focus on detecting image tampering by analyzing inconsistencies in pixel patterns, compression artifacts, and statistical features. However, with the rapid advancement of image editing tools and artificial intelligence, attackers are increasingly developing anti-forensic techniques designed to hide traces of manipulation and evade forensic detection systems. This creates significant challenges for reliable forgery localization. This study proposes a novel anti-forensics framework aimed at bypassing digital forgery localization methods by intelligently modifying manipulated regions to mimic authentic image characteristics. The framework utilizes advanced image processing and machine learning techniques to alter forensic traces such as noise patterns, compression artifacts, and edge inconsistencies that are typically used by forensic detection algorithms. By integrating adaptive feature manipulation and learning-based optimization, the proposed system reduces the effectiveness of existing forgery localization models. The framework evaluates its performance against common image forensic detection methods and demonstrates its capability to obscure manipulation evidence while maintaining high visual quality of the altered images. Experimental results show that the proposed approach significantly decreases the detection accuracy of traditional forensic models while preserving perceptual realism. The research highlights the growing need for more robust and resilient digital forensic systems capable of handling advanced anti-forensic attacks. The proposed framework contributes to understanding the limitations of current forensic techniques and encourages the development of stronger detection mechanisms for securing digital media authenticity.

INTRODUCTION

In the modern digital era, images play a crucial role in communication, journalism, social media, legal investigations, and scientific documentation. The rapid advancement of digital imaging technologies and image editing software has made it easier to manipulate digital images with high precision. While these technologies offer significant benefits for creative and professional purposes, they also create opportunities for malicious manipulation of visual content. Image forgeries can be used to spread misinformation, manipulate public opinion, fabricate evidence, or damage reputations. As a result, ensuring the authenticity and integrity of digital images has become an important challenge in the field of digital forensics.

Digital image forensics focuses on detecting and analyzing manipulated images to determine whether an image has been altered. One of the major research areas in digital forensics is **forgery localization**, which aims to identify the exact regions within an image that have been manipulated. Common types of image forgeries include copy-move forgery, splicing, retouching, and object removal. To detect these manipulations, researchers have developed several forensic techniques based on statistical

analysis, signal processing, and machine learning. These techniques analyze inconsistencies in image properties such as noise patterns, color distribution, compression artifacts, and structural features to detect signs of manipulation.

In recent years, **deep learning and artificial intelligence (AI)** have significantly improved the accuracy of forgery detection and localization systems. Convolutional Neural Networks (CNNs), transformer-based models, and other deep learning architectures have been widely used to identify manipulated regions in images with high precision. These models learn complex patterns and hidden features within images that are difficult to detect using traditional forensic techniques. As a result, modern forgery localization systems can effectively identify various forms of image manipulation even when the alterations are subtle or sophisticated.

However, as forensic detection methods become more advanced, **anti-forensics techniques** have also evolved to counter these detection mechanisms. Anti-forensics refers to methods used to conceal evidence of manipulation or to mislead forensic analysis tools. The goal of anti-forensics is to modify digital images in such a way that forensic detectors fail to identify the manipulated regions. For example, anti-forensic techniques may remove traces of editing operations, adjust image noise patterns, manipulate compression artifacts, or introduce synthetic textures that mimic natural image characteristics. By altering these features, attackers attempt to make forged images appear authentic to both human observers and automated detection systems.

The development of anti-forensic techniques poses a significant challenge for digital forensics researchers. As detection methods improve, attackers continuously adapt their strategies to bypass these systems. This ongoing competition between forensic detection and anti-forensic evasion creates a dynamic research landscape where new methods are constantly required to stay ahead of emerging threats. In particular, advanced image editing tools and generative models such as deepfakes and diffusion-based image generation systems have increased the complexity of detecting manipulated images. These technologies allow attackers to generate highly realistic synthetic images that can be difficult to distinguish from genuine ones.

To address these challenges, researchers are exploring **novel frameworks that analyze the interaction between forgery localization methods and anti-forensic strategies**. Instead of focusing solely on detection, these frameworks examine how manipulated images can evade forensic analysis and how detection systems can be strengthened against such evasion attempts. Understanding the mechanisms used in anti-forensics helps researchers design more robust detection algorithms and improve the reliability of digital forensic tools.

The proposed study introduces a **novel framework for anti-forensics in digital forgery localization** that aims to investigate and model techniques used to evade detection systems. The framework focuses on understanding how forged images can be



modified to avoid localization algorithms while maintaining visual realism. By analyzing different manipulation strategies and their impact on forensic detection models, the framework provides insights into the vulnerabilities of current detection techniques.

In this framework, advanced image processing methods and machine learning techniques are utilized to simulate anti-forensic operations and evaluate their effectiveness against existing forgery localization systems. The framework may include modules for image manipulation, feature alteration, and evaluation of detection performance. By studying the relationship between manipulation strategies and detection accuracy, researchers can identify weaknesses in current forensic algorithms and develop improved countermeasures.

Furthermore, the integration of artificial intelligence within the proposed framework enables automated analysis of large image datasets and detection models. AI-based approaches can learn patterns associated with successful anti-forensic attacks and help identify defensive strategies that enhance the resilience of forensic systems. This approach not only contributes to a deeper understanding of anti-forensic techniques but also supports the development of more secure and reliable image authentication technologies.

The importance of this research extends beyond academic interest. Digital images are widely used as evidence in legal proceedings, journalism, security investigations, and social media platforms. If forged images can successfully evade detection systems, the consequences may include misinformation, legal disputes, and threats to public trust. Therefore, improving the robustness of digital forensic technologies is essential for maintaining the credibility of visual information in the digital age.

In conclusion, the rapid growth of image manipulation technologies has increased the need for advanced digital forensic solutions capable of detecting sophisticated forgeries. At the same time, the emergence of anti-forensic techniques highlights the need for continuous research into methods that can counter these evasion strategies. The proposed novel framework for anti-forensics in digital forgery localization aims to analyze and understand the mechanisms used to bypass forensic detection systems while providing insights that can lead to the development of more robust and secure image authentication technologies. Through the integration of image processing, machine learning, and forensic analysis, this research contributes to strengthening the reliability and effectiveness of digital image forensics in combating emerging threats in the digital media landscape.

LITERATURE REVIEW

Digital images play a critical role in modern communication, journalism, legal investigations, and social media platforms. However, the rapid advancement of image editing software and artificial intelligence has made it easier to manipulate images. As a result, **digital image forensics** has emerged as an important research area that focuses on detecting image tampering and identifying manipulated regions. In contrast, **anti-forensics** aims to hide or alter the traces that forensic techniques rely on, making it difficult to detect image manipulation or forgery. This literature review examines existing research on digital forgery localization and anti-forensic techniques, highlighting the progress, challenges, and research gaps in this field.

Early research in digital image forensics mainly focused on detecting image manipulation using statistical and signal

processing techniques. One of the most widely studied approaches is **Error Level Analysis (ELA)**, which detects inconsistencies in compression levels within an image. Since digital images are often compressed using JPEG, manipulated regions may have different compression artifacts compared to the original image. Researchers have used ELA and similar methods to identify tampered areas in images. However, these methods are sensitive to image quality and compression levels, making them vulnerable to anti-forensic attacks that modify compression artifacts.

Another important approach used in digital image forensics is **sensor pattern noise analysis**. Each digital camera sensor produces a unique noise pattern that can act as a fingerprint for identifying the source of an image. Researchers have used this technique to detect tampered regions by analyzing inconsistencies in sensor noise patterns. If a part of an image has been manipulated or copied from another source, the noise pattern may not match the rest of the image. Although this method has shown promising results, anti-forensic techniques have been developed to remove or alter sensor noise patterns, making it more difficult for forensic tools to detect manipulation.

With the advancement of machine learning and deep learning technologies, researchers have introduced **data-driven approaches** for detecting digital forgeries. Convolutional Neural Networks (CNNs) have been widely used for forgery detection and localization tasks. These models can automatically learn complex patterns in images and identify subtle inconsistencies caused by editing operations such as copy-move, splicing, or retouching. Deep learning models such as ResNet, VGGNet, and U-Net have demonstrated high accuracy in detecting manipulated regions in images. Despite these advancements, deep learning-based forensic models can still be vulnerable to adversarial attacks and anti-forensic techniques designed to confuse or mislead the detection algorithms.

Anti-forensics has become an emerging research area that focuses on **developing methods to hide evidence of image manipulation**. One common anti-forensic technique involves modifying JPEG compression artifacts to make manipulated regions appear consistent with the rest of the image. By adjusting compression parameters or re-compressing the entire image, attackers can reduce the effectiveness of compression-based forensic methods. Another approach involves adding artificial noise to the image in order to mimic the sensor noise pattern of a specific camera, thereby deceiving sensor-based forensic analysis.

In recent years, researchers have explored the use of **generative models and adversarial learning** for anti-forensic applications. Generative Adversarial Networks (GANs) can be used to create highly realistic images and manipulate image content while maintaining visual consistency. GAN-based techniques can also be used to remove forensic traces from manipulated images, making it difficult for traditional detection algorithms to identify tampered regions. For example, some studies have demonstrated that GANs can be trained to modify images in a way that preserves visual quality while eliminating statistical artifacts used by forensic detectors.

Another important development in anti-forensics is the use of **adversarial examples** to fool deep learning-based forensic models. Adversarial attacks involve making small and often imperceptible changes to an image that cause machine learning models to produce incorrect predictions. In the context of digital



forensics, adversarial perturbations can be applied to manipulated images so that forensic algorithms fail to detect the forged regions. These techniques highlight the need for more robust forensic models that can resist adversarial manipulation.

Researchers have also investigated **hybrid anti-forensic frameworks** that combine multiple techniques to evade detection. For instance, some frameworks integrate noise manipulation, compression artifact modification, and deep learning-based image enhancement to conceal forgery traces. Such hybrid approaches can significantly reduce the effectiveness of existing forensic tools, especially when the manipulated images are shared on social media platforms where images are frequently compressed or resized.

Despite significant progress in both forensic detection and anti-forensic methods, several challenges remain in this field. One major challenge is the **arms race between forensic analysts and attackers**. As forensic techniques become more sophisticated, attackers develop new anti-forensic methods to bypass detection. Another challenge is the rapid growth of **AI-generated content**, including deepfake images and synthetic media. These technologies make it increasingly difficult to distinguish between authentic and manipulated content.

Furthermore, many existing forensic models rely on specific assumptions about image acquisition or compression, which may not hold true for all images. Variations in camera devices, image editing tools, and compression algorithms can affect the reliability of forensic analysis. Therefore, researchers are exploring more generalized and robust approaches that can detect manipulation across diverse datasets and imaging conditions.

In summary, digital image forgery detection has evolved from traditional statistical analysis to advanced deep learning-based methods capable of identifying complex manipulation patterns. At the same time, anti-forensic techniques have also advanced significantly, leveraging machine learning, generative models, and adversarial attacks to hide traces of image manipulation. The development of a **novel anti-forensic framework for digital forgery localization** aims to explore new strategies for evading forensic detection while maintaining high visual quality in manipulated images. Future research should focus on improving the robustness of forensic algorithms, developing datasets for evaluating anti-forensic techniques, and designing secure systems capable of resisting adversarial manipulation in digital media environments.

SYSTEM ANALYSIS

EXISTING SYSTEM

Current digital image forensics systems are designed to **detect and localize image forgeries** such as copy-move manipulation, splicing, and image editing. These systems use techniques such as **statistical feature analysis, machine learning algorithms, and deep learning models** to identify inconsistencies in digital images. They analyze artifacts like noise patterns, compression traces, color inconsistencies, and pixel-level irregularities to detect manipulated regions.

Common approaches include **Convolutional Neural Networks (CNNs), error level analysis (ELA), and feature-based forgery detection methods**. These techniques help investigators and security agencies identify altered images and maintain digital authenticity.

However, modern image editing tools and sophisticated manipulation techniques can sometimes bypass these detection systems. As a result, researchers have started exploring **anti-forensics techniques**, which intentionally modify images to hide traces of manipulation and evade detection systems.

Disadvantages of Existing System

- **Limited robustness against advanced anti-forensic attacks**
- Difficulty detecting **highly sophisticated image manipulations**
- **High computational complexity** in deep learning-based detection systems
- Vulnerability to **adversarial image modifications**
- Reduced accuracy when images undergo **multiple editing operations**
- Difficulty distinguishing between **natural image artifacts and manipulated traces**

PROPOSED SYSTEM

The proposed system introduces a **novel anti-forensics framework designed to evade digital forgery localization systems** by intelligently modifying manipulated images to hide forensic traces. This framework focuses on **altering statistical patterns and image features** that are typically used by forensic detection algorithms.

The system uses **advanced machine learning and image processing techniques** to analyze how forgery detection models identify manipulated regions. Based on this understanding, the framework applies targeted modifications such as **noise pattern adjustment, compression artifact manipulation, and feature-level obfuscation** to reduce the likelihood of detection.

The proposed framework may also incorporate **adversarial learning techniques**, where the anti-forensic model learns to generate manipulated images that appear natural and evade detection models. This approach enables the system to **adapt to different forensic detection methods** and improve the stealthiness of manipulated content.

ADVANTAGES OF PROPOSED SYSTEM

- Improved ability to **bypass modern forgery detection systems**
- **Adaptive anti-forensic techniques** that evolve with detection models
- Enhanced ability to **hide manipulation traces in digital images**
- Increased robustness against **deep learning-based forensic detectors**
- Better understanding of **weaknesses in existing forensic detection systems**
- Supports research in developing **more secure and robust forensic detection frameworks**



IMPLEMENTATION

1. Data Collection Module

The implementation begins with collecting **digital image datasets containing both authentic and forged images**. These datasets may include manipulated images created using techniques such as **copy-move forgery, image splicing, retouching, and object removal**. Public datasets like CASIA, Columbia, or other forensic datasets can be used to train and evaluate the system.

2. Data Preprocessing Module

In this stage, the collected images are prepared for analysis. Preprocessing includes **resizing images, noise filtering, color normalization, and image format standardization**. This step ensures that the images are consistent and suitable for further feature extraction and model training.

3. Forgery Feature Extraction Module

This module extracts **image forensic features** that are commonly used for detecting image manipulation. Examples of extracted features include:

- Pixel inconsistencies
- Compression artifacts
- Edge and texture irregularities
- Color distribution patterns

Techniques such as **Discrete Cosine Transform (DCT), wavelet transforms, and texture analysis** can be applied to identify potential forgery traces.

4. Forgery Localization Module

In this phase, machine learning or deep learning models such as **Convolutional Neural Networks (CNNs)** are trained to detect and localize manipulated regions in images. The model analyzes spatial patterns and identifies specific regions where tampering has occurred. The output is usually a **forgery localization map** highlighting suspicious areas.

5. Anti-Forensics Generation Module

This module focuses on developing **anti-forensic techniques** to hide traces of digital manipulation. Various methods are applied to modify image characteristics so that forensic detection algorithms cannot easily identify forged regions. Techniques may include:

- Noise pattern modification
- Compression artifact adjustment
- Texture smoothing and blending
- Metadata alteration

These methods attempt to make manipulated images appear natural and indistinguishable from authentic images.

6. Adversarial Learning Module

An **adversarial learning framework** can be implemented where an anti-forensic model attempts to evade detection while a forensic detection model attempts to identify manipulation. This interaction helps improve both detection and anti-forensic strategies, often using **Generative Adversarial Networks (GANs)**.

7. Model Evaluation Module

The system performance is evaluated using metrics such as:

- **Detection accuracy**
- **Precision and recall**
- **F1-score**
- **Localization accuracy**

These metrics help measure how effectively the anti-forensic techniques reduce the detectability of forged regions.

8. Visualization and Analysis Module

Visualization tools are used to compare **original images, forged images, and anti-forensically modified images**. Heatmaps or segmentation masks may be used to highlight detected forgery regions and evaluate the effectiveness of the anti-forensic techniques.

9. System Deployment Module

The final framework can be deployed as a **research or forensic evaluation tool** to test the robustness of digital image forensic systems. It can help researchers understand the limitations of existing detection methods and develop stronger forensic techniques.

10. Algorithm

DECISION TREE CLASSIFIERS

Decision tree classifiers are used successfully in many diverse areas. Their most important feature is the capability of capturing descriptive decision making knowledge from the supplied data. Decision tree can be generated from training sets. The procedure for such generation based on the set of objects (S), each belonging to one of the classes $C_1, C_2, \dots, C_k$ is as follows:

Step 1. If all the objects in S belong to the same class, for example C_i , the decision tree for S consists of a leaf labeled with this class

Step 2. Otherwise, let T be some test with possible outcomes $O_1, O_2, \dots, O_n$. Each object in S has one outcome for T so the test partitions S into subsets $S_1, S_2, \dots, S_n$ where each object in S_i has outcome O_i for T. T becomes the root of the decision tree and for each outcome O_i we build a subsidiary decision tree by invoking the same procedure recursively on the set S_i .

GRADIENT BOOSTING Gradient boosting is a [machine learning](#) technique used in [regression](#) and [classification](#) tasks, among others. It gives a prediction model in the form of an [ensemble](#) of weak prediction models, which are typically [decision trees](#).^{[1][2]} When a decision tree is the weak learner, the resulting algorithm is called gradient-boosted trees; it



usually outperforms [random forest](#). A gradient-boosted trees model is built in a stage-wise fashion as in other [boosting](#) methods, but it generalizes the other methods by allowing optimization of an arbitrary [differentiable loss function](#).

K-NEAREST NEIGHBORS (KNN)

- Simple, but a very powerful classification algorithm
- Classifies based on a similarity measure
- Non-parametric
- Lazy learning
- Does not “learn” until the test example is given
- Whenever we have a new data to classify, we find its K-nearest neighbors from the training data

Example

- Training dataset consists of k-closest examples in feature space
- Feature space means, space with categorization variables (non-metric variables)
- Learning based on instances, and thus also works lazily because instance close to the input vector for test or prediction may take time to occur in the training dataset

LOGISTIC REGRESSION CLASSIFIERS

Logistic regression analysis studies the association between a categorical dependent variable and a set of independent (explanatory) variables. The name *logistic regression* is used when the dependent variable has only two values, such as 0 and 1 or Yes and No. The name *multinomial logistic regression* is usually reserved for the case when the dependent variable has three or more unique values, such as Married, Single, Divorced, or Widowed. Although the type of data used for the dependent variable is different from that of multiple regression, the practical use of the procedure is similar. Logistic regression competes with discriminant analysis as a method for analyzing categorical-response variables. Many statisticians feel that logistic regression is more versatile and better suited for modeling most situations than is discriminant analysis. This is because logistic regression does not assume that the independent variables are normally distributed, as discriminant analysis does. This program computes binary logistic regression and multinomial logistic regression on both numeric and categorical independent variables. It reports on the regression equation as well as the goodness of fit, odds ratios, confidence limits, likelihood, and deviance. It performs a comprehensive residual analysis including diagnostic residual

reports and plots. It can perform an independent variable subset selection search, looking for the best regression model with the fewest independent variables. It provides confidence intervals on predicted values and provides ROC curves to help determine the best cutoff point for classification. It allows you to validate your results by automatically classifying rows that are not used during the analysis.

NAÏVE BAYES

The naive bayes approach is a supervised learning method which is based on a simplistic hypothesis: it assumes that the presence (or absence) of a particular feature of a class is unrelated to the presence (or absence) of any other feature. Yet, despite this, it appears robust and efficient. Its performance is comparable to other supervised learning techniques. Various reasons have been advanced in the literature. In this tutorial, we highlight an explanation based on the representation bias. The naive bayes classifier is a linear classifier, as well as linear discriminant analysis, logistic regression or linear SVM (support vector machine). The difference lies on the method of estimating the parameters of the classifier (the learning bias). While the Naive Bayes classifier is widely used in the research world, it is not widespread among practitioners which want to obtain usable results. On the one hand, the researchers found especially it is very easy to program and implement it, its parameters are easy to estimate, learning is very fast even on very large databases, its accuracy is reasonably good in comparison to the other approaches. On the other hand, the final users do not obtain a model easy to interpret and deploy, they does not understand the interest of such a technique. Thus, we introduce in a new presentation of the results of the learning process. The classifier is easier to understand, and its deployment is also made easier. In the first part of this tutorial, we present some theoretical aspects of the naive bayes classifier. Then, we implement the approach on a dataset with Tanagra. We compare the obtained results (the parameters of the model) to those obtained with other linear approaches such as the logistic regression, the linear discriminant analysis and the linear SVM. We note that the results are highly consistent. This largely explains the good performance of the method in comparison to others. In the second part, we use various tools on the same dataset ([Weka 3.6.0](#), [R 2.9.2](#), [Knnime 2.1.1](#), [Orange 2.0b](#) and [RapidMiner 4.6.0](#)). We try above all to understand the obtained results.

RANDOM FOREST

Random forests or random decision forests are an ensemble learning method for classification, regression and other tasks that operates by constructing a multitude of decision trees at training time. For classification tasks, the output of the random forest is the class selected by most trees. For regression tasks, the mean or average prediction of the individual trees is returned. Random decision forests correct for decision trees' habit of overfitting to their training set. Random forests generally outperform decision trees, but their accuracy is lower than gradient boosted trees.



However, data characteristics can affect their performance. The first algorithm for random decision forests was created in 1995 by Tin Kam Ho[1] using the random subspace method, which, in Ho's formulation, is a way to implement the "stochastic discrimination" approach to classification proposed by Eugene Kleinberg. An extension of the algorithm was developed by Leo Breiman and Adele Cutler, who registered "Random Forests" as a trademark in 2006 (as of 2019, owned by Minitab, Inc.). The extension combines Breiman's "bagging" idea and random selection of features, introduced first by Ho[1] and later independently by Amit and Geman[13] in order to construct a collection of decision trees with controlled variance. Random forests are frequently used as "blackbox" models in businesses, as they generate reasonable predictions across a wide range of data while requiring little configuration.

SVM

In classification tasks a discriminant machine learning technique aims at finding, based on an *independent and identically distributed (iid)* training dataset, a discriminant function that can correctly predict labels for newly acquired instances. Unlike generative machine learning approaches, which require computations of conditional probability distributions, a discriminant classification function takes a data point x and assigns it to one of the different classes that are a part of the classification task. Less powerful than generative approaches, which are mostly used when prediction involves outlier detection, discriminant approaches require fewer computational resources and less training data, especially for a multidimensional feature space and when only posterior probabilities are needed. From a geometric perspective, learning a classifier is equivalent to finding the equation for a multidimensional surface that best separates the different classes in the feature space. SVM is a discriminant technique, and, because it solves the convex optimization problem analytically, it always returns the same optimal hyperplane parameter—in contrast to *genetic algorithms (GAs)* or *perceptrons*, both of which are widely used for classification in machine learning. For perceptrons, solutions are highly dependent on the initialization and termination criteria. For a specific kernel that transforms the data from the input space to the feature space, training returns uniquely defined SVM model parameters for a given training set, whereas the perceptron and GA classifier models are different each time training is initialized. The aim of GAs and perceptrons is only to minimize error during training, which will translate into several hyperplanes' meeting this requirement.

CONCLUSION

In this paper, we introduce the novel task of pixel-level anti forensics and propose a unified framework, SEAR, to combat forgery localization algorithms and evade detection at the pixel level. SEAR employs a concealer to generate perturbation maps and add them to original forged images, producing anti-forensic images that obscure tampering traces and deceive forgery

localization models. By incorporating self-supervised training, the concealer learns to constrain disturbances to tampered regions and remove high-frequency traces, resulting in visually imperceptible effects on other regions. Additionally, adversarial learning integrates a well-trained forgery localization network into the system, and after a combination of self-supervised and adversarial training, the concealer can slightly modify original forged images and successfully deceive forgery localization models. We provide detailed experimental results to validate SEAR. This work paves the way for further exploration of anti-forensics. Interestingly, adversarial learning benefits both anti-forensic algorithms and forensic models and could inspire future work on forgery localization. Our future research focuses on improving the robustness of forensic algorithms based on anti-forensic methods.

REFERENCES

- [1] L. Verdoliva, "Media forensics and deepfakes: an overview," *IEEE Journal of Selected Topics in Signal Processing*, vol. 14, no. 5, pp. 910–932, 2020.
- [2] S. Lyu, X. Pan, and X. Zhang, "Exposing region splicing forgeries with blind local noise estimation," *International Journal of Computer Vision*, vol. 110, no. 2, pp. 202–221, 2014.
- [3] D. Cozzolino, G. Poggi, and L. Verdoliva, "Efficient dense-field copy–move forgery detection," *IEEE Transactions on Information Forensics and Security*, vol. 10, no. 11, pp. 2284–2297, 2015.
- [4] Y. Fan, P. Carr'e, and C. Fernandez-Maloigne, "Image splicing detection with local illumination estimation," in *Proceedings of International Conference on Image Processing*. IEEE, 2015, pp. 2940–2944.
- [5] J. G. Han, T. H. Park, Y. H. Moon, and I. K. Eom, "Efficient Markov feature extraction method for image splicing detection using maximization and threshold expansion," *Journal of Electronic Imaging*, vol. 25, no. 2, p. 023031, 2016.
- [6] C. Li, Q. Ma, L. Xiao, M. Li, and A. Zhang, "Image splicing detection based on Markov features in QDCT domain," *Neurocomputing*, vol. 228, pp. 29–36, 2017.
- [7] P. Zhou, X. Han, V. I. Morariu, and L. S. Davis, "Learning rich features for image manipulation detection," in *Proceedings of Conference on Computer Vision and Pattern Recognition*. IEEE, 2018, pp. 1053–1061.
- [8] J. H. Bappy, C. Simons, L. Nataraj, B. Manjunath, and A. K. Roy-Chowdhury, "Hybrid LSTM and encoder–decoder architecture for detection of image forgeries," *IEEE Transactions on Image Processing*, vol. 28, no. 7, pp. 3286–3300, 2019.
- [9] Y. Wu, W. AbdAlmageed, and P. Natarajan, "Mantra-net: Manipulation tracing network for detection and localization of image forgeries with anomalous features," in *Proceedings of*



Conference on Computer Vision and Pattern Recognition. IEEE, 2019, pp. 9543–9552.

[10] X. Hu, Z. Zhang, Z. Jiang, S. Chaudhuri, Z. Yang, and R. Nevatia, “SPAN: spatial pyramid attention network for image manipulation localization,” in Proceedings of European Conference on Computer Vision. Springer, 2020, pp. 312–328.

[11] L. Zhuo, S. Tan, B. Li, and J. Huang, “Self-adversarial training incorporating forgery attention for image forgery localization,” IEEE

Transactions on Information Forensics and Security, vol. 17, pp. 819–834, 2022.

[12] X. Guo, X. Liu, Z. Ren, S. Grosz, I. Masi, and X. Liu, “Hierarchical fine-grained image forgery detection and localization,” in Proceedings

of Conference on Computer Vision and Pattern Recognition, 2023, pp. 3155–3165.

[13] F. Guillaro, D. Cozzolino, A. Sud, N. Dufour, and L. Verdoliva, “Trufor: Leveraging all-round clues for trustworthy image forgery detection and localization,” in Proceedings of Conference on Computer Vision and Pattern Recognition, 2023, pp. 20 606–20 615.

[14] M. Kirchner and R. Bohme, “Hiding traces of resampling in digital images,” IEEE Transactions on Information Forensics and Security,

vol. 3, no. 4, pp. 582–592, 2008.

[15] C.-W. Kwok, O. C. Au, and S.-H. Chui, “Alternative anti-forensics method for contrast enhancement,” in Proceedings of International Workshop on Digital Watermarking. Springer, 2011, pp. 398–410.