
Splicing Tampering Detection Algorithm Design for Digital Media Image Privacy

Yuxuan Liu* and Siyi Feng

Academy of Fine Arts, NanJing XiaoZhuang University, 210000 Nanjing, China
E-mail: liuyx1232026@outlook.com

**Corresponding Author*

Received 26 January 2026; Accepted 13 April 2026

Abstract

To make image information more authentic and complete and to prevent splicing tampered image data from affecting social fairness, this paper proposes an algorithm based on Deep Convolutional Neural Networks to detect splicing tampered digital media image privacy. Building upon deep convolutional feature extraction, this algorithm introduces a self-attention mechanism to enhance focus on tampered regions. For the first time, it innovatively applies a boundary-aware loss function to patch tampering detection, effectively addressing the issue of ambiguous boundary region detection and significantly improving localization accuracy. The experiment was conducted on datasets from the Institute of Automation, Chinese Academy of Sciences, Cover Dataset, National Institute of Standards and Technology datasets, and the 2020 Image Tampering Dataset. The algorithm demonstrated the following performance metrics: area under the receiver operating characteristic curve values of 0.971, 0.961, and 0.987, respectively; accuracy rate of 98.94%; precision rate of 97.12%; recall rate of 99.16%; F1 mean values of 0.941 and 0.952 under gamma ray and noise interference, respectively, indicating strong robustness. These results prove that the proposed

Journal of Cyber Security and Mobility, Vol. 15.4, 777–798.

doi: 10.13052/jcsm2245-1439.1541

© 2026 River Publishers

algorithm can achieve precise detection of splicing tampered image privacy. It effectively addresses the problem of insufficient detection accuracy in some existing methods. It also promotes the intelligent development of detection and contributes to building a more authentic information environment in society.

Keywords: Digital media image privacy, splicing tampering, deep convolutional neural network, attention mechanism, boundary-aware loss function.

1 Background

With the rapid development of digital media, image communication has become a key carrier of social interaction and human information transmission [1, 2]. Image editing tools are evolving toward more convenient and intelligent functions, which makes splicing tampering images easier to create [3, 4]. Splicing tampering, as a common way of combining different images to change the original meaning of a picture, can conceal or falsify information [5]. It involves user privacy and can be exploited in online fraud, fake news, and other improper situations [6, 7]. These issues seriously damage the authenticity and trust of digital media and modern society [8, 9]. Designing effective algorithms for image splicing tampering detection helps protect personal privacy and network security [10, 11]. Current mainstream splicing tampering detection methods primarily include traditional approaches based on manually designed features and deep learning models. Traditional methods rely on manually designed features, resulting in limited generalization capabilities. While existing deep learning methods have improved detection accuracy, they still exhibit shortcomings in boundary region localization and robustness against complex post-processing interference [12, 13]. Therefore, enhancing model boundary perception capabilities and anti-interference performance while improving detection accuracy remains a critical research gap in current studies. Deep Convolutional Neural Networks (DCNNs) have the advantages of extracting deep feature information and strong robustness. They can improve detection accuracy and reduce diagnostic time [14]. Therefore, this study proposes an algorithm based on DCNNs for detecting mosaic tampering in digital media images, aiming to enhance detection accuracy and localization clarity. This method addresses the two key challenges of fuzzy boundary region localization and reduced robustness under complex

interference conditions in existing models. It is expected to efficiently identify mosaic-tampered images, thereby providing reliable technical support for image authenticity verification and privacy protection. The key innovations of this algorithm include: introducing a self-attention mechanism into DCNNs to enhance global perception of tampered regions and incorporating boundary-aware loss functions. These advancements provide novel technical approaches for tampering recognition in digital media image splicing applications.

2 Literature Review

As a powerful algorithm for feature extraction, DCNN improved overall recognition accuracy when applied to different fields. Many scholars have carried out studies in this direction. Pande et al. raised a model based on three-dimensional DCNN for liver image separation. The preprocessed liver image data were input into the model for feature extraction. The iterative region growing technique improved prediction accuracy and achieved image segmentation [15]. Siale et al. proposed an algorithm based on DCNN for large-scale abnormal monitoring data. The algorithm extracted features and combined synthetic minority oversampling techniques for data enhancement and integration, which effectively solved abnormal detection problems. The results showed that the algorithm reached a training accuracy of 98.47% on an Internet of Things threat network dataset [16]. To enhance the accuracy of feature extraction, Dhamale and his team proposed a framework based on parallel DCNN. The framework enhanced image data and information extraction by minimizing interference in images and highlighting object features [17]. Adigopula et al. proposed a model for authentication in radio frequency fingerprint systems. The model input raw samples and processed data transmitted them through different models and quickly identified label attributes. The results showed that the model achieved 98% label recognition accuracy [18]. Facing the challenge that a single method was difficult to extract emotion information, Ye and Xiao proposed a model combining DCNN and multimodal sentiment analysis. The model introduced selective kernel networks and bidirectional long short-term memory networks to strengthen feature extraction. It searched modality interactions through CNN and attention mechanisms. The results showed that the model reached 98.00% accuracy on a Kaggle text dataset [19].

As a medium for information transmission, digital media images played an important role in ensuring authenticity. Many scholars domestically and internationally have carried out studies on splicing tampering detection. Xing et al. proposed a framework based on dual-channel enhanced attention dense convolutional networks for tampering detection in power images. The framework improved detection accuracy through the backbone network, linked global context information through the encoder structure, and optimized parameters through feature maps. The results showed that the framework improved evaluation indicators by 30% [20]. Shi et al. proposed a lightweight local tampering detection method based on convolutional network MobileNetV2 and dual stream network. This method improves MobileNetV2 by retaining richer image tampering traces, and then combines dual stream networks to extract image tampering features and noise features of real regions [21]. Ding and other scholars proposed a new method for image tampering localization based on dual channel U-Net. This method uses an encoder, feature fusion, and decoder for image detection, and constructs a dual channel encoding network model to analyze the original tampered image and tampered residual image. The fused feature map is input into the decoder to decode the predicted image layer by layer [22]. Researchers such as Alsughayer proposed a detection model based on the U-net architecture to accurately identify traces of tampering in remote sensing images. This model extracts residual noise, introduces a constrained convolutional layer to locate tampering in the case of concatenation, and finally trains it using a conditional generative adversarial network framework [23]. Zhang and other collaborators proposed a dual branch tampering image detection method based on multi-scale features for image tampering detection. This method introduces a fusion module with attention mechanism to improve the sensitivity of the network to tampered areas, and then uses the rich edge information of shallow features as guidance to identify subtle differences between tampered and unaltered areas. The experimental results show that the detection F1 score of this method is 0.766 [24].

In summary, current studies have achieved progress in splicing tampering detection. However, recognition and prediction accuracy still need further improvement. DCNNs have the advantage of autonomous feature learning, which improved detection accuracy of tampered images. Therefore, this paper proposes a DCNN-based algorithm for detecting splicing tampered digital media image privacy. The goal was to prevent distortion and miscommunication of real information and to contribute to social privacy protection.

3 Algorithm Design for Splicing Tampered Digital Media Image Privacy

3.1 Splicing Tampering Detection Algorithm Design Based on DCNN

Because DCNN extracts target features from large-scale data at multiple levels, applying it to splicing tampering detection of digital media image privacy can handle complex post-processing such as Gaussian blur and noise addition. It also identifies higher-level problems caused by splicing tampering [25, 26]. Therefore, this paper proposes an algorithm based on DCNNs to detect splicing tampered images. The algorithm extracts features of digital media image privacy through the encoder-decoder structure of DCNN. It collects the feature size of splicing tampering on images and outputs a probability map of tampering detection. DCNNs consist of four main structures, which are the convolutional layer, the pooling layer, the activation layer, and the normalization layer. Its structural diagram is shown in Figure 1.

As shown in Figure 1, DCNN first preprocesses input image data by normalization and other operations to prepare for detection. The processed data enter the convolutional layer for analysis. The convolutional layer extracts features through sliding convolution kernels and calculates weighted sums of local features to generate a new feature map. The pooling layer samples the feature map to reduce computation, performing pooling on each region to obtain a smaller feature map. The activation layer introduces nonlinear functions into the feature map, which enables the algorithm to learn more complex content. The normalization layer performs batch normalization and dropout operations to prevent overfitting and to improve training speed. The algorithm then predicts a probability map of splicing tampering and outputs it as the result. The feature map expression of the convolutional layer is shown

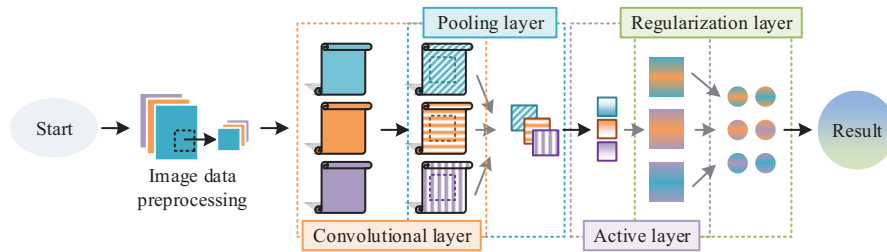


Figure 1 Schematic diagram of the structure of DCNN.

in Equation (1).

$$f_j^{(n)} = \sum_i \omega_{ij}^{(n)} \cdot f_j^{(n-1)} + b_j^{(n)} \quad (1)$$

where $f_j^{(n)}$ represents the output feature map of the j -th channel in the n -th layer, $b_j^{(n)}$ represents the bias, and $\omega_{ij}^{(n)}$ represents the convolution kernel weight. Data of the convolutional layer are input as multiple feature maps. The pooling layer performs pooling on data and produces a smaller feature map, which is expressed in Equation (2).

$$f_j^{(n+1)} = \text{pool}(f_j^{(n)}) \quad (2)$$

where $f_j^{(n+1)}$ represents the feature map of the pooling layer, and pool represents the pooling function. The activation layer transforms the feature map through nonlinear functions, and the new feature map expression is shown in Equation (3).

$$f_j'^{(n+1)} = \text{Activation}(f_j^{(n)}) \quad (3)$$

where $f_j'^{(n+1)}$ represents the feature map of the activation layer, and Activation represents the activation function. The normalization layer produces a normalized feature map, which is expressed in Equation (4).

$$f(x_i) = \frac{x_i - m}{v + \text{eps}} \cdot \gamma + \beta \quad (4)$$

where γ and β represent learning parameters, m represents the mean value of input data, eps is a constant to avoid division by zero, and v is the variance from input data. Normalization reduces the risk of gradient vanishing and improves convergence speed. Self-Attention (SA) reduces the interference of background and environmental factors, which makes the algorithm focus on splicing tampered regions of image privacy [27]. Based on DCNNs, this study adopts SA to optimize the structure and forms SA-DCNN to achieve more precise and efficient detection of splicing tampered digital media image privacy. SA strengthens feature representation through query, key, and value vectors, and its expression is shown in Equation (5).

$$\begin{cases} Q = XW^Q \\ K = XW^K \\ V = XW^V \end{cases} \quad (5)$$

where W^Q represents the weight matrix of query vectors, W^K represents the weight matrix of key vectors, and W^V represents the weight matrix of value vectors. The dot product between key and query vectors is calculated by Equation (6).

$$A_{ij} = \frac{Q_i K_j^T}{\sqrt{d_k}} \quad (6)$$

where A_{ij} represents the correlation between feature i and feature j , and $\sqrt{d_k}$ represents the dimension of key vectors. Attention weights are normalized through a SoftMax function, and the expression is shown in Equation (7).

$$Attention(Q, K, V) = \text{softmax}(A)V \quad (7)$$

where softmax represents the SoftMax function. Multi-head attention splits features into multiple subspaces, calculates attention separately, and outputs the final splicing result. The computation of each head is shown in Equation (8).

$$head = Attention(QW_i^Q, KW_i^K, VW_i^V) \quad (8)$$

The attention mechanism enables the algorithm to capture dependencies between features in subspaces and to build deeper understanding of feature sequences. The structural diagram of SA is shown in Figure 2.

As shown in Figure 2, the algorithm inputs image privacy data and extracts features. The image features are linearly transformed into query, key, and value vectors. Then the algorithm calculates the dot product of query and key vectors, scales and normalizes it, and obtains a probability distribution. This distribution builds feature correlations and multiplies with value vectors. Weighted sums generate output features, which produce a new sequence with global context information, and the algorithm outputs the final result. The computation of multi-head attention is shown in Equation (9).

$$MultiHead(Q, K, V) = \text{Concat}(head_1, head_2, \dots, head_n)W^o \quad (9)$$

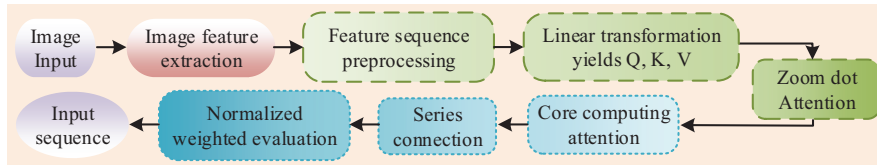


Figure 2 Schematic diagram of the structure of SA.

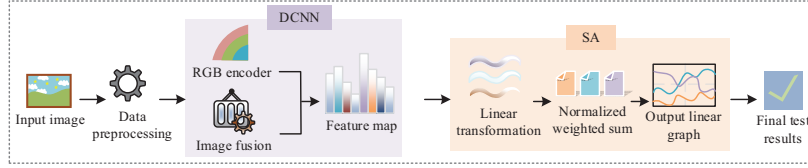


Figure 3 Operation process of SA-DCNN hybrid algorithm.

where *Concat* represents the splicing operation, W^o represents the output projection weight matrix, and n represents the number of attention heads. SA-DCNN integrates local feature extraction and global information interaction, which enables more precise detection of splicing tampered digital media image privacy. Its structural diagram is shown in Figure 3.

In Figure 3, the algorithm first inputs images collected from digital media and preprocesses privacy data. In the DCNN module, the encoder, which is composed of convolutional layers, activation functions, and pooling layers, extracts local textures and features. Different feature maps are fused and input into the next module. The SA module performs linear transformations on feature maps and calculates weighted sums of query, key, and value vectors. It captures dependencies between information and outputs an enhanced feature map with global context information. The final feature map represents the detected splicing tampered region of the image.

3.2 Optimization Design of Detection Algorithm for Splicing Tampered Images

Although SA-DCNN can effectively detect the splicing tampered parts of image privacy, it still needs further improvement in generalization ability and resource cost. The Boundary-Aware Loss Function (BAL), because of its boundary-centered feature, clearly calculates the spatial region between the predicted boundary and the ground truth label. It makes the algorithm positioning clearer and improves convergence [28]. Therefore, this paper applies the optimized BAL-SA-DCNN algorithm to splicing tampering detection of image privacy to achieve high accuracy, high efficiency, and strong generalization in identification and prediction. The probability expression of BAL for probability map prediction is shown in Equation (10).

$$P_{softmax} = Softmax(P) \quad (10)$$

where P is the probability map that satisfies $P \in R^{B \times C \times D \times H \times W}$. $P_{softmax}$ is the predicted probability. The maximum probability is taken as the result

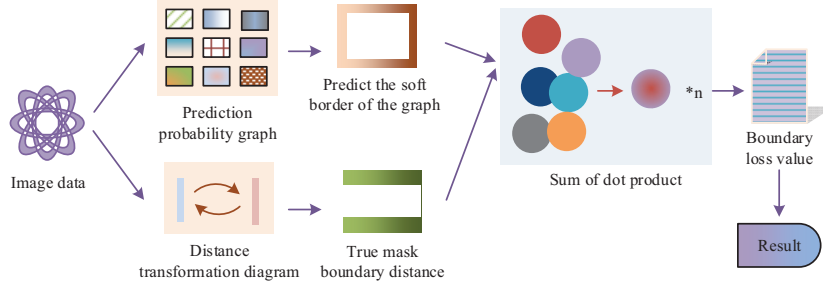


Figure 4 Schematic diagram of the structure of BAL.

of the predicted segmentation, as shown in Equation (11).

$$S_{pred} = \arg \max_c P_{softmax}[:, c, :, :] \quad (11)$$

where S_{pred} is the predicted segmentation result. The calculation expression of the ground truth label and the predicted segmentation result is shown in Equation (12).

$$\begin{cases} B_{gt} = DetectBoundary(S_{gt}) \\ B_{pred} = DetectBoundary(S_{pred}) \end{cases} \quad (12)$$

where S_{gt} is the ground truth label. S_{pred} is the predicted segmentation result. B_{gt} and B_{pred} are the boundaries of S_{gt} and S_{pred} . The core region mask obtained by integrating the boundary map is shown in Equation (13).

$$M_{critical} = B_{gt} \vee B_{pred} \quad (13)$$

where $M_{critical}$ is the key region mask. $\vee$ represents the logical “or” operation. The structure of BAL is shown in Figure 4.

As shown in Figure 4, the algorithm first converts the input image data into the ground truth mask. Then the ground truth mask is processed through two steps simultaneously. In the first step, the distance transform of the ground truth mask is calculated, generating a distance transform map of the same size. In the second step, the soft boundary of the predicted map is calculated to obtain the region reflecting the predicted boundary. The two steps are then combined by dot product and summation, and the algorithm finally outputs external pixels of the ground truth boundary and internal pixels of the non-ground truth boundary. This produces the final boundary loss and detects the splicing tampered region of image privacy. The process of calculating standard cross-entropy loss is shown in Equation (14).

$$L_{CE}(P, S_{gt})[b, 0, d, h, w] = -\log(P_{softmax}[b, 0, d, h, w], d, h, w) \quad (14)$$

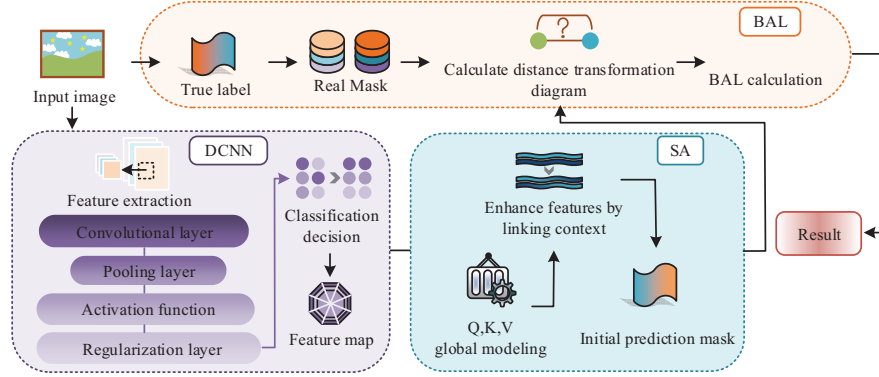


Figure 5 Recognition process of the BAL-SA-DCNN algorithm.

where L_{CE} is a tensor with the same size as the input data. The cross-entropy loss value of each data point is stored in the corresponding position for subsequent weighting. The process of BAL weighting the loss by the ground truth mask is shown in Equation (15).

$$L_{boundary} = \frac{\sum_{b,d,h,w} L_{CE}(P, S_{gt})[b, 0, d, h, w] \cdot M_{critical}[b, 0, d, h, w]}{\sum_{b,d,h,w} M_{critical}[b, 0, d, h, w] + \varepsilon} \quad (15)$$

where ε is a small constant to prevent the denominator from being zero. The optimized BAL-SA-DCNN algorithm has efficient computational performance and accurately identifies and judges the splicing tampered parts of images. Its structure is shown in Figure 5.

As shown in Figure 5, the program first inputs the image privacy information and performs preprocessing, then enters the DCNN module. If it is a supervision signal or ground truth label, it directly enters the BAL module for processing. In the BAL module, the ground truth label is converted into a ground truth mask, and the distance transform map is calculated to obtain the distance loss value, which is used to output the prediction result. In the DCNN module, data go through convolutional layers, activation functions, and pooling layers to extract features, which are then input into the SA module in different feature map forms. The SA module builds a model of the feature map data through query vector, key vector, and value vector. After linear transformation, the model produces the predicted mask, which enters the BAL module for processing. The ground truth mask is processed by BAL, and the final result is obtained.

4 Performance Evaluation of Splicing Tampering Detection Algorithm Based on DCNN

4.1 Performance Verification of Splicing Tampering Detection by SA-DCNN

To verify the superiority of SA-DCNN in detecting the splicing tampered parts of image privacy, the study compared it with Manipulation Tracing Network (Mantra), Semantic Predictions Aided Network (SPAN), and Multi-View Stream Structure Network (MVSSN). The experimental system parameters were Windows 10, deep learning framework PyTorch, programming language Python 3.8, NVIDIA GeForce RTX4060 Laptop, 16GB RAM, learning rate 0.0001, and 100 iterations. The datasets included the Institute of Automation, Chinese Academy of Sciences (CASIA) dataset, the Copy-move forgery coverage dataset (Coverage), the National Institute of Standards and Technology (NIST) dataset, the Image Manipulation Detection 2020 dataset (IMD2020), and a self-made dataset. The content of the self-made dataset is shown in Table 1.

As shown in Table 1, the self-made dataset collected images with tampering operations such as vertical and horizontal flipping, image compression, and Gaussian noise addition. It also randomly sampled from CASIA and NIST datasets, with a total of 10,650 images. The study tested SA-DCNN, Mantra, SPAN, and MVSSN on CASIA, Coverage, and NIST datasets for image classification accuracy. The area under the receiver operating

Table 1 Details of the self-made dataset

Type	Image Type	Number of Images
Test set	Original image	300
	Splicing tampered images	250
	JPEG compression	1000
	Gaussian noise	1000
	CASIA	400
	NIST	600
Training set	Original image	600
	Splicing tampered images	500
	JPEG compression	1200
	Gaussian noise	1800
	CASIA	1200
	NIST	1800

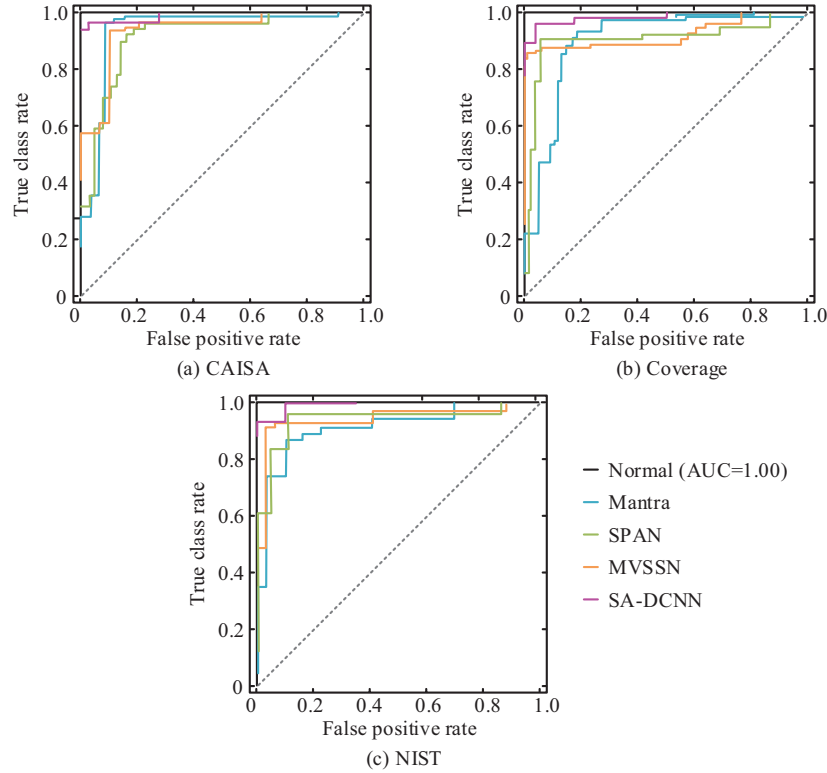


Figure 6 Comparison of AUC test results.

characteristic curve (AUC) was used to evaluate the detection performance. The results are shown in Figure 6.

As shown in Figure 6, the AUC value represented the ability to distinguish authentic images from tampered or spliced images. In Figures 6(a–c), SA-DCNN achieved AUC values closer to those of authentic images, with values of 0.971, 0.961, and 0.987. The AUC values of the other three methods were all lower. In summary, the closer the AUC value of SA-DCNN was to 1, the better classification performance it showed. It demonstrated stronger capability to distinguish authentic images from spliced images, effectively detecting authentic pixels and reducing the risk of misguidance by tampered images, thereby providing more reliable data. To further verify the reliability of SA-DCNN, the study tested the four methods on CASIA and NIST datasets, using accuracy, precision, and recall as indicators. The results are shown in Figure 7.

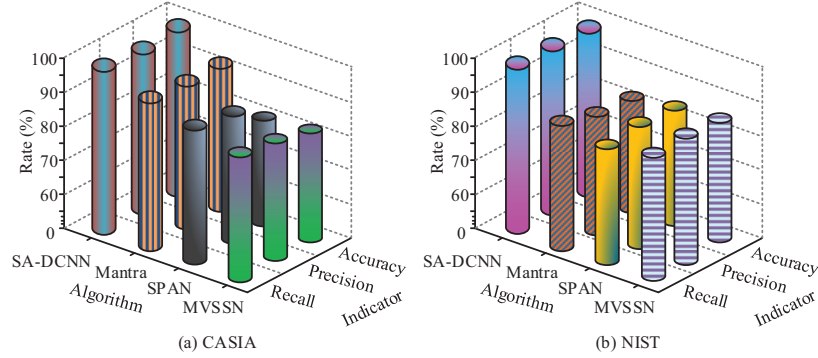


Figure 7 Comparison of test results of accuracy, precision, and recall.

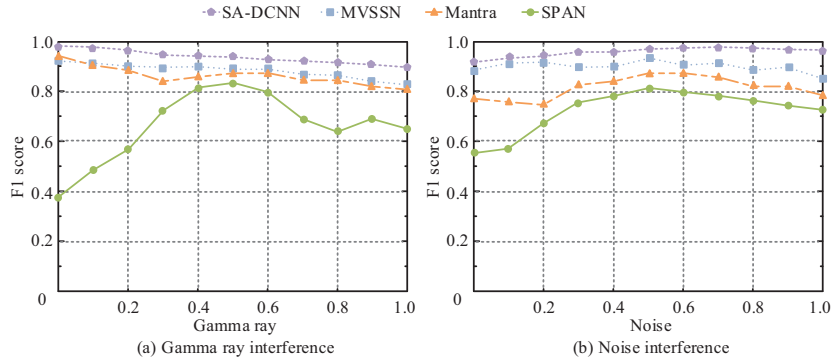


Figure 8 Comparison of F1 score test results.

In Figure 7, SA-DCNN achieved higher values in accuracy, precision, and recall compared with the other methods. On the CASIA dataset, SA-DCNN reached an accuracy of 98.94%, precision of 97.12%, and recall of 99.16%. On the NIST dataset, it reached an accuracy of 99.16%, precision of 99.13%, and recall of 99.11%. Among the methods, SPAN performed the worst. These results indicated that SA-DCNN achieved high prediction accuracy with very few false detections, effectively controlled missed detections, and showed better overall performance than the other methods. It reduced errors and provided more reliable and comprehensive data, ensuring the authenticity of image privacy. To further confirm its superiority, the study tested the four methods on the IMD2020 dataset for robustness, simulating changes in brightness and contrast by varying gamma intensity and adding noise interference. The results are shown in Figure 8.

As shown in Figure 8(a), as gamma intensity increased, the F1 score of SA-DCNN decreased slowly, with an average of approximately 0.941. In Figure 8(b), the F1 score of SA-DCNN increased slightly with more noise, with an average of approximately 0.952. The other three methods had lower F1 scores with greater fluctuations. These results indicated that SA-DCNN achieved higher F1 scores under different types and intensities of interference, showing stronger robustness. It was less sensitive to complex external disturbances and was not easily misled when identifying and detecting spliced images, proving its stability and robustness.

4.2 Effect Verification of Splicing Tampering Detection Based on DCNN

After verifying the performance of SA-DCNN in splicing detection, the study further evaluated BAL-SA-DCNN. The study compared it with Mantra, SPAN, and MVSSN. The experimental system was Windows 10 with PyTorch, Python 3.8, NVIDIA GeForce RTX4060 Laptop, 16GB RAM, learning rate 0.0001, and 100 iterations. The datasets included CASIA, Coverage, NIST, IMD2020, and the self-made dataset. The study trained and validated the four methods on CASIA, Coverage, NIST, and IMD2020 datasets, using cross-entropy loss as the indicator to evaluate generalization performance. The results are shown in Figure 9.

As shown in Figure 9(a), BAL-SA-DCNN achieved cross-entropy loss values of 0.015, 0.010, 0.012, and 0.011 on the training sets of CASIA, Coverage, NIST, and IMD2020. In Figure 9(b), its cross-entropy loss values on the validation sets of these datasets were 0.014, 0.011, 0.016, and 0.009.

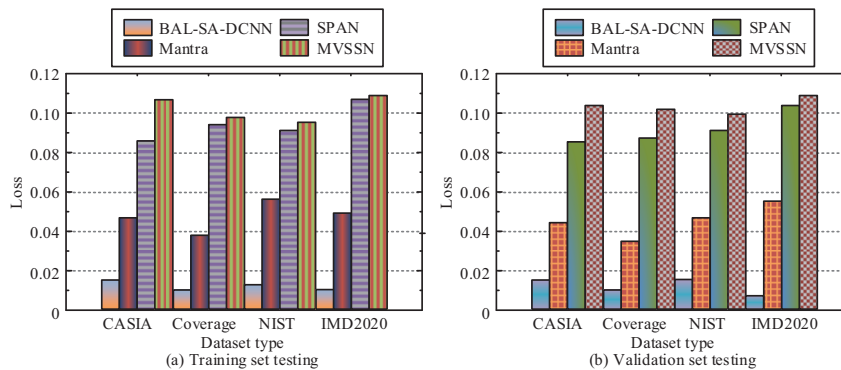


Figure 9 Comparison of cross entropy loss test results.

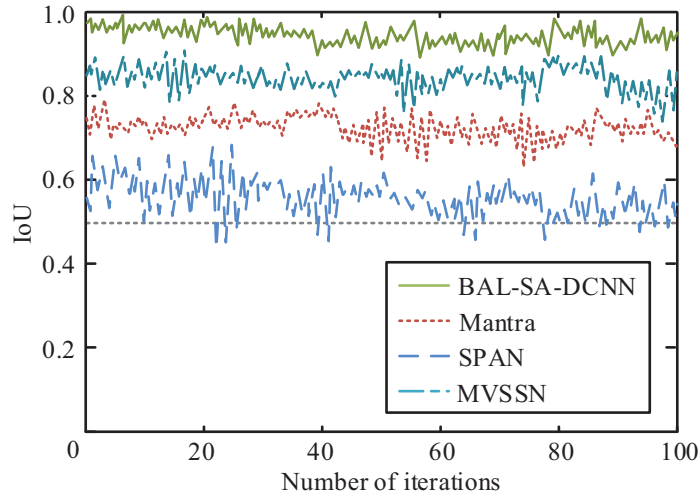


Figure 10 Comparison of IoU results of each algorithm's prediction results.

Moreover, Figure 9 showed that MVSSN performed the worst, with loss values on CASIA increasing by 0.094 and 0.095 compared with the proposed algorithm. These results showed that BAL-SA-DCNN achieved lower loss values on both training and validation sets. Because the BAL mechanism learned patterns during training, it fit the data well, achieved clear boundary separation, improved prediction of splicing tampering in image privacy, and produced highly reliable results, proving its strong generalization ability. To further evaluate detection performance, the study tested the four methods on IMD2020, using Intersection over Union (IoU) to evaluate prediction accuracy. The results are shown in Figure 10.

As shown in Figure 10, a higher IoU indicated a higher probability of correctly predicting tampered regions. BAL-SA-DCNN achieved an average IoU of about 0.954, higher than the other three methods. SPAN performed the worst with an average IoU of about 0.563. These results showed that BAL-SA-DCNN had stronger ability to understand the shape, size, and boundary of tampered regions, achieving higher positioning accuracy. The predicted regions highly matched the ground truth, helping provide reliable data for image privacy. The self-made dataset included 200 face replacement images, 200 background replacement images, 300 object addition or removal images, 300 copy-move images, 350 semantic inconsistency splicing images, 350 text or logo tampering images, 480 certificate tampering images, 420 medical image tampering images, 310 news tampering images, and 290

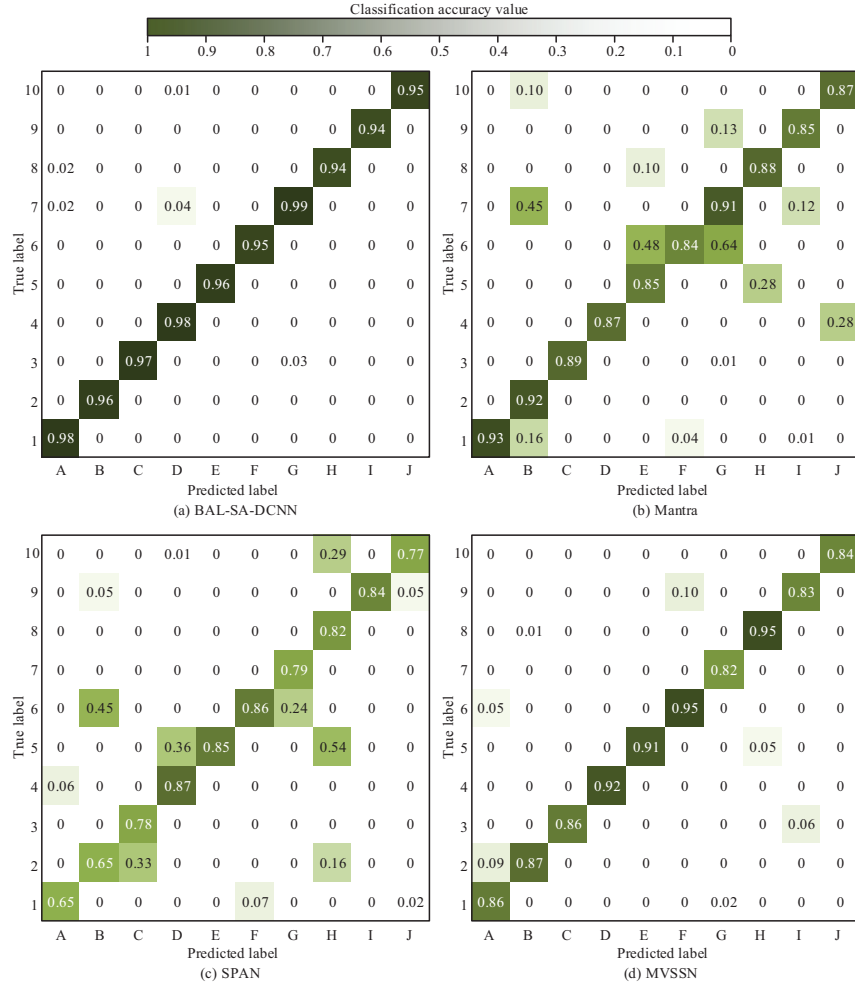


Figure 11 Confusion matrix of image type classification accuracy.

academic misconduct tampering images, totaling 3200 images. These types were labeled A to J. To further analyze the hybrid algorithm, the study tested the four methods on the self-made dataset for classification accuracy. The results are shown in Figure 11.

As shown in Figure 11(a), BAL-SA-DCNN achieved an average classification accuracy of 0.96, while SPAN performed the worst with an accuracy of 0.79. These results demonstrated that BAL-SA-DCNN performed significantly better in classifying spliced or tampered images. Its attention

mechanism focused more on tampered regions, and its feature extraction capability improved classification accuracy. It was a reliable algorithm that effectively and accurately detected splicing and tampering in digital media image privacy.

5 Summary and Future Work

Because spliced and tampered images might distort the truth and threaten the privacy of digital media images, and because the existing detection methods still had problems such as insufficient accuracy and low efficiency, this study proposes an algorithm based on DCNNs for detecting splicing tampering of image privacy. Under boundary constraints, this algorithm uses the feature extraction ability of DCNNs combined with the attention mechanism to identify and classify spliced and tampered images, which achieved efficient and accurate detection and protected the privacy of digital media images. The experimental results showed that the SA-DCNN algorithm achieved AUC values of 0.971, 0.961, and 0.987 on three datasets, with an accuracy of 98.94%, a precision of 97.12%, and a recall of 99.16%. The average F1 values under gamma ray and noise interference were 0.941 and 0.952, which showed the strongest robustness. The BAL-SA-DCNN algorithm achieved cross-entropy loss values of 0.015, 0.010, 0.012, and 0.011 on the training sets of CASIA, Coverage, NIST, and IMD2020 datasets, with an average IoU of approximately 0.954 and an average classification accuracy of 0.96. In summary, the DCNNs-based algorithm showed excellent accuracy, precision, classification performance, and robustness in detecting spliced and tampered images, and its performance could meet the current needs for image privacy detection. For instance, in the field of digital forensics, research algorithms can assist in identifying tampered content within legal cases and news reports. For social media platforms, these algorithms automatically detect and flag spliced images to curb the spread of misinformation. By enabling more accurate and efficient image authentication, such research algorithms help preserve information integrity and enhance privacy protection within digital ecosystems. Although the proposed algorithm demonstrates superior performance in feature extraction, classification, prediction, and recognition, its detection speed remains underexplored. Future research should focus on three key directions: first, enhancing computational efficiency through techniques like algorithm pruning and knowledge distillation; second, evaluating the algorithm across diverse datasets to further validate its generalization capabilities; third, improving real-time detection performance to enable rapid

processing of high-resolution images for real-time monitoring applications. Future in-depth studies are anticipated to optimize algorithmic performance across all dimensions, expand its applicability to various fields, protect digital media image privacy, and provide society with safer, more reliable data support.

Fundings

The research is supported by The Jiangsu Provincial College Students' Innovation and Entrepreneurship Training Program project titled "Development of Digital Cultural and Creative Products for Museums Based on AI Technology," guided in December 2024. Project number: 202411460102Y; The 2025 General Research Project on Teaching Reform in Higher Education in Jiangsu Province titled "Ideological and Political Nurturing, Digital Empowerment, and Aesthetic Education Practice: Research on the Construction of a 'Three-Dimensional Synergistic' Aesthetic Education System in Normal Colleges." Project number: 2025JGYB473; The 2025 National-level Project of the Jiangsu Provincial College Students' Innovation and Entrepreneurship Training Program titled "Breakthrough and Rebirth: AI Digital and Intelligent Inheritance and Innovation Research on the Intangible Cultural Heritage of Zisha Pottery by 'Jiping Zaowu'."

References

- [1] Jędrasiak K, Wolański R. An image Analysis Algorithm for Detecting Modified Content on the Internet Against Cybercrime. A Technique for Estimating the Probability of Modification. *Safety & Fire Technology*, 2024, 63(1): 88–94.
- [2] Bai X, Bai Y. Equilibrium Strategy of Attack and Defense in Computer Networks Based on Markov Signal Game Theory. *Journal of Cyber Security and Mobility*, 2025, 14(1): 127–154. DOI:0.13052/jcsm2245-1439.1416.
- [3] Hasanvand M, Nooshyar M, Moharamkhani E, Selyari A. Machine Learning Methodology for Identifying Vehicles Using Image Processing. *AIA*, 2023, 3(1): 170–178.
- [4] Somsuk K. Enhanced Algorithm for Recovering RSA Plaintext when Two Modulus Values Share at least One Common Prime Factor. *Journal of Cyber Security and Mobility*, 2025, 14(2): 433–456. DOI:10.13052/jcsm2245-1439.1427.

- [5] Dutta M, Saini J. A systematic literature review on image splicing detection and localization using emerging technologies. *Multimedia Tools and Applications*, 2025, 84(19): 20721–20756.
- [6] Li X, Zhang J. Multi-source Data Fusion for Real-time Cybersecurity Situational Awareness and Visualization. *Journal of Cyber Security and Mobility*, 2025, 14(4): 955–980. DOI:10.13052/jcsm2245-1439.1448.
- [7] Xing J, Tian X, Han Y. A Dual-channel Augmented Attentive Dense-convolutional Network for power image splicing tamper detection. *Neural Computing and Applications*, 2024, 36(15): 8301–8316.
- [8] Kumari R, Garg H. Image splicing forgery detection: A review. *Multimedia Tools and Applications*, 2025, 84(8): 4163–4201.
- [9] Akram A, Jaffar M A, Rashid J, Boulaaras S M, Faheem M. CMV2U-net: au-shaped network with edge-weighted features for detecting and localizing image splicing. *Journal of Forensic Sciences*, 2025, 70(3): 1026–1043.
- [10] Hu J, Xue R, Teng G, et al. Image splicing manipulation location by multi-scale dual-channel supervision. *Multimedia Tools and Applications*, 2024, 83(11): 31759–31782.
- [11] Chen E. Analysis of E-commerce Security Protection Technology Based on YOLO Algorithm Optimized by Lightweight Neural Network. *Journal of Cyber Security and Mobility*, 2025, 14(4): 849–876. DOI:10.13052/jcsm2245-1439.1444.
- [12] Kumari R, Garg H. Image splicing forgery detection: A review. *Multimedia Tools and Applications*, 2025, 84(8): 4163–4201.
- [13] Dutta M, Saini J. A systematic literature review on image splicing detection and localization using emerging technologies. *Multimedia Tools and Applications*, 2025, 84(19): 20721–20756.
- [14] Subha S, Kumaran U. Efficient Liver Segmentation using Advanced 3D-DCNN Algorithm on CT Images. *Engineering, Technology & Applied Science Research*, 2025, 15(1): 19324–19330.
- [15] Pande S D, Kalyani P, Nagendram S, Alluhaidan A S, Badu G H, Ahammad S H, Pandey V K, Kumar A, Sridevi G, Bonyah E. Comparative analysis of the DCNN and HFCNN Based Computerized detection of liver cancer. *BMC Medical Imaging*, 2025, 25(1): 37–50.
- [16] Siale A Y D, Hassan Q M Z, Kadekle M F A S, Veena B S. Enhancing Large-Scale Network Security with a VGG-Net-Based DCNN: A Deep Learning Approach to Anomaly Detection. *Journal of Robotics and Control (JRC)*, 2025, 6(3): 1316–1331.

- [17] Adigopula S, Subramanyam M V. DCNN-RFF-NFC: a novel design of NFC security using deep Convolution neural network-based RF fingerprinting. *Neural Computing and Applications*, 2025, 37(6): 4439–4453.
- [18] Dhamale T, Bhandari S, Harpale V, Sakhi P, Napte K, Mahajan A. Autism spectrum disorder detection using parallel DCNN with improved teaching learning optimization feature selection scheme. *SAIEE Africa Research Journal*, 2025, 116(3): 89–100.
- [19] Ye H, Xiao X. MEA-IFE: An Improved Multi-modal Fusion Framework Based on DCNN-BERT-BiLSTM and Its Application in Sentiment Analysis. *Information Technology and Control*, 2025, 54(2): 451–470.
- [20] Xing J, Tian X, Han Y. A Dual-channel Augmented Attentive Dense-convolutional Network for power image splicing tamper detection. *Neural Computing and Applications*, 2024, 36(15): 8301–8316.
- [21] Shi X, Li P, Wu H, Chen Q, Zhu H. A lightweight image splicing tampering localization method based on MobileNetV2 and SRM. *IET Image Processing*, 2023, 17(6): 1883–1892.
- [22] Ding H, Chen L, Tao Q, Fu Z, Dong L, Cui X. DCU-Net: a dual-channel U-shaped network for image splicing forgery detection. *Neural computing and applications*, 2023, 35(7): 5015–5031.
- [23] Alsughayer R, Hussain M, Saeed F, AboalSamh H. Detection and localization of splicing on remote sensing images using image-to-image transformation. *Applied Intelligence*, 2023, 53(11): 13275–13292.
- [24] Zhang H, Guo C, Wang X. Double-branch forgery image detection based on multi-scale feature fusion. *Optoelectronics Letters*, 2024, 20(5): 307–312.
- [25] Saleh A, Mahmoud K, Hussein M M, Nasrat L, Nasser M A. Efficient 2D DCNN approach for detecting and classifying faults in modular power converters. *Journal of Integrated Science and Technology*, 2025, 13(6): 1137.
- [26] Waghumbare A, Singh U. DIAT-DSCNN-GRU-HARNet: A Lightweight DCNN for Video Based Classification of Human Activities. *Automatic Control and Computer Sciences*, 2025, 59(2): 255–265.
- [27] Wang J, Shao H, Peng Y, Liu B. PSparseFormer: Enhancing fault feature extraction based on parallel sparse self-attention and multiscale broadcast feedforward block. *IEEE Internet of Things Journal*, 2024, 11(13): 22982–22991.
- [28] Li C, Zhang J, Niu D, Zhao X, Yang B, Zhang C. Boundary-aware uncertainty suppression for semi-supervised medical image segmentation. *IEEE Transactions on Artificial Intelligence*, 2024, 5(8): 4074–4086.

Biographies



Yuxuan Liu obtained her Ph.D. in Design from Dongmyeong University, South Korea, in 2020. She is currently working as an Associate Professor and Deputy Dean of the School of Fine Arts at Nanjing Xiaozhuang University. She has been deeply engaged in the field of digital media art for many years, publishing more than 20 academic papers and receiving over 30 awards and honors.



Siyi Feng obtained her Bachelor of Arts degree in Visual Communication Design from Nanjing Xiaozhuang University in 2025. She currently serves as the General Manager of Nanjing Rangchen Cultural Media Co. Ltd. She independently founded the original Zisha cultural and creative brand Jiping, which successfully secured RMB 300,000 in seed funding invested by Nanjing Zijing Venture Capital Co. Ltd. She possesses outstanding capabilities in overall brand operation and aesthetic design. Specializing in visual design, she is proficient in packaging design and integrated brand VI design, and has received numerous recognitions in the cultural and creative industry.

