

Consistency-Heterogeneity Balanced Fake News Detection via Cross-Modal Matching

Ying Guo , Bingxin Li , Kexin Zhen , Jie Liu , Gaolei Li , *Member, IEEE*, Qi Wang , *Member, IEEE*, and Yong-Jin Liu , *Senior Member, IEEE*

Abstract—Generating synthetic content through generative AI (GAI) presents considerable hurdles for current fake news detection methodologies. Many existing detection approaches concentrate on feature-based multimodal fusion, neglecting semantic relationships such as correlations and diversities. In this study, we introduce an innovative cross-modal matching-driven approach to reconcile semantic relevance (text–image consistency) and semantic gap (text–image heterogeneity) in multimodal fake news detection. Unlike the conventional paradigm of multimodal fusion followed by detection, our approach integrates textual modality, visual modality (images), and text embedded within images (auxiliary modality) to construct an end-to-end framework. This framework considers the relevance of contents across different modalities while simultaneously addressing the gap in structures, achieving a delicate balance between consistency and heterogeneity. Consistency is fostered by evaluating intermodality correlation via pairwise-similarity scores, while heterogeneity is addressed by employing cross-attention mechanisms to account for intermodality diversity. To achieve equilibrium between consistency and heterogeneity, we employ attention-guided enhanced modality interaction and similarity-based dynamic weight assignment to establish robust frameworks. Comparative experiments conducted on the Chinese Weibo dataset and the English Twitter dataset demonstrate the effectiveness of our approach, surpassing the state-of-the-art by 7% to 13%.

Impact Statement—Fake news detection is a popular technology in social media’s security. It helps to clean social contents and create harmonious environments to a large degree by means of artificial intelligence. However, recent research has demonstrated

that two paradoxical concepts, i.e. text–image consistency and text–image heterogeneity contribute to the authenticity of a news together. The consistency-heterogeneity balanced algorithm we introduce in this article overcame these limitations. With a significant increase in detection accuracy by 7% to 13% on Chinese Weibo dataset and the English Twitter dataset, the technology is ready to support users in a wide variety of applications including governmental decisions, industrial designs, and individual overcomes.

Index Terms—Cross-modal matching, multimodal fake news detection, text–image consistency, text–image heterogeneity.

I. INTRODUCTION

GENERATIVE artificial intelligence (GAI), including diffusion models, deepfakes, and ChatGPT, has emerged as a foundational element within social networks. The rapid progression of GAI has led to a significant surge in the proliferation of fake news. This deliberately fabricated information has inflicted considerable harm upon the public, underscoring the importance of implementing measures to combat this issue and ensure the dissemination of reliable and accurate information. Fake news has garnered widespread attention in recent years due to its profound negative impact on significant public events, such as the U.S. presidential election, the Russia–Ukraine conflict, the COVID-19 pandemic, and others. It is worth emphasizing that rumors spread on social media may impact the opinions of billions of people [1]. Consequently, the detection of fake news assumes paramount importance and serves as the primary focus of this article.

Current techniques for detecting fake news can be broadly categorized into two groups: single-modal based [2], [3], [4] and multimodal based [5], [6], [7]. Given that instances of fake news often encompass a blend of diverse data types—such as user comments, manipulated images, fabricated videos, and other variables—detecting multimodal fake news remains an intricate challenge. To date, diverse multimodal approaches have concentrated on distinct modality fusion techniques to detect fake news [8], [9], [10]. For instance, Wang et al. [8] utilized a straightforward splicing approach to amalgamate features from various modalities, whereas others had enhanced fusion capabilities by employing strategies such as attention [9] and bilinear pooling [10].

Nowadays, semantics are proving to provide evident clues for identifying fake news. In the incorporation of semantic features, the SAFE framework [11] was the first to introduce image–text consistency, while the MCNN framework [12] introduced

Received 29 April 2024; revised 28 October 2024; accepted 3 January 2025.
Date of publication 13 January 2025; date of current version 30 June 2025.
This work was supported in part by the Ministry of Education in China (MOE) Liberal Arts and Social Sciences Foundation under Grant 24YJC86007, in part by the Key Supported Program of Joint Fund of the National Natural Science Foundation of China under Grant U23B2029, in part by the National Natural Science Foundation of China under Grant 62076167, and in part by Yuxiu Innovation Project of NCUT under Grant 2024NCUTYXCX102. This article was recommended for publication by Associate Editor Chaker Abdelaziz Kerrache upon evaluation of the reviewers’ comments. (*Corresponding author: Ying Guo.*)

Ying Guo, Bingxin Li, Kexin Zhen, and Jie Liu are with the North China University of Technology, Beijing 100144, P. R. China (e-mail: guoying@ncut.edu.cn; libingxin1126@163.com; kexinzen@163.com; liujxxy@126.com).

Gaolei Li is with Shanghai Jiao Tong University, Shanghai 200240, China (e-mail: gaolei_li@sjtu.edu.cn).

Qi Wang is with the State Key Laboratory of Public Big Data, Guizhou University, Guizhou 550025, China (e-mail: qiwang@gzu.edu.cn).

Yongjin Liu is with the Department of Computer Science and Technology, Tsinghua University, Beijing 100084, China (e-mail: liuyongjin@tsinghua.edu.cn).

Digital Object Identifier 10.1109/TAI.2025.3527921

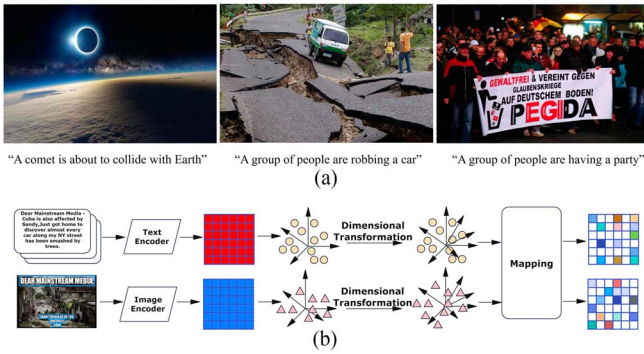


Fig. 1. Concerns regarding text-image consistency and heterogeneity. (a) Incongruity between the accompanying text and the image is readily apparent, giving rise to suspicions about the news being characterized as misinformation. (b) Notable semantic disparity is present between the text and the image, potentially affecting the overall comprehensibility.

image tampering to propose a refined model based on it. Despite these breakthroughs, the text-image consistency is limited to the main modalities and neglects the textual parts embedded in visual elements, which may also be useful [13]. On the other hand, there exist complex mapping relationships in terms of “whether text and image are consistent or not” and “whether the news is true or not”. Evidently, the discord between the semantics of text and image stands as a pivotal indicator of fabricated news, as depicted in Fig. 1(a). But as illustrated in Fig. 1(b), even when dealing with accurate news, a semantic gap can emerge due to the distinct modalities. Therefore, the text-image consistency generally reclines to the content-related consistency and neglects the structure-related inconsistency. That is, the text-image heterogeneity, which is an intrinsic but losing aspect, goes hand in hand with the text-image consistency. It is desired to solve the issues below since consistency and heterogeneity are contradictory concepts: 1) How to cooperate between the contents of texts and images, concerning the semantic relevance; 2) How to debate between the structures of texts and images, concerning the semantic gap; and 3) How to reconcile these two conflicting concepts to attain equilibrium.

In this article, our objective is to strike a delicate balance between text-image consistency and text-image heterogeneity. To achieve this, we propose an end-to-end multimodal fake news detection framework based on cross-modal matching (MFND-CMM), comprising four key components: *Modality integration*: We employ a novel multimodal framework to extract textual clues from attached images, enriching the content with more diverse and comprehensive information. *Feature extraction*: We utilize an improved Bert encoder, an enhanced ResNet encoder, and a refined OCR encoder for extracting textual, visual, and auxiliary features, respectively. *Feature interaction*: To facilitate interaction among different features, we employ cross-attention mechanisms, addressing semantic heterogeneity arising from intermodality diversity across different modalities. *Feature matching*: In this component, we assess the semantic consistency among different modalities by evaluating intermodality correlation through pairwise-similarity scores.

Building upon this foundation, this end-to-end network adeptly achieves a balance between consistency and heterogeneity, by employing attention-guided enhanced modality interaction and similarity-based dynamic weight assignment. These techniques contribute to the development of robust frameworks designed to effectively predict the news labels. The primary contributions of this article are outlined as follows.

- 1) We propose a novel multimodal fake news detection framework driven by cross-modal matching. The model effectively extracts fine-grained salient features of news, promoting semantic consistency among different modalities while alleviating semantic heterogeneity.
- 2) We consider the text embedded within the image as one of the crucial determinants to integrate a multimodal encoder. On the basis of it, we emphasize intermodality correlation by pairwise-similarity scores, retain intermodality diversity by cross-attention mechanism, and achieve a trade-off between consistency and heterogeneity via attention-guided enhanced modality interaction and similarity-based dynamic weight assignment.
- 3) We empirically show that the proposed model can effectively identify fake news and outperform the state-of-the-art multimodal fake news detection models on two large-scale real-world datasets, especially on English Twitter.

II. RELATED WORK

A. Fake News Detection

Both language-based and vision-based fake news detection have received widespread attention. For the language-based methods, Castilo et al. [2] utilized some word-level metrics such as the number of words and the number of links to infer fake news. Qazvinian et al. [4] used Bayesian network as a classifier to detect fake news and mainly studied the topic features of the text. This method of manually extracting features has major limitations. It is not only time-consuming and labor-intensive, but also cannot capture complex feature relationships. Ma et al. [3] used recurrent neural networks to extract implicit text information from texts and pioneered a rumor text feature extraction method based on deep learning. Qian et al. [14] developed a text-based method to extract semantic information from article text and proposed a user-response generation model to assist in fake news detection. Bhattarai et al. [15] used the vocabulary and semantic attributes of text to detect fake news. And for the vision-based methods, Gupta et al. [16] conducted preliminary work based on images, and for the first time established a classification model for false image recognition on Twitter by artificially extracting relevant information such as the time when false images appeared on Twitter. However, this method cannot obtain the semantic features of the image. Thanks to the development of deep learning, many studies use convolutional neural networks for image feature extraction. Qi et al. [17] proposed to identify fake news by extracting image spatial and frequency domain features to determine image tampering. Cao et al. [18] argued that typical image manipulation detection methods are useful in revealing traces of news tampering. Zhang

et al. [19] used pretrained convolutional neural networks to model fake news images and extracted deep visual features.

However, for multimodal news, these methods fail to detect cross-modal correlations (although these unimodal features can be explored, and they do play a key role in distinguishing fake news, correlations and consistency are ignored and other multimodal features, which may harm the overall performance of these single-modal schemes on multimodal news). The key issue in multimodal fake news detection is adapting linguistic and visual representations.

B. Multimodal Fusion

In recent years, some new fusion techniques have been introduced into fake news detection, replacing the simple concatenation. Wang et al. [8] proposed EANN, which uses an auxiliary task of event classification to improve versatility. Singhal et al. [20] proposed Spotfake, which uses VGG and BERT to extract image and text features respectively, and connect them for classification. However, since these two methods only rely on fusion features obtained directly using connection or attention mechanisms, the text and image features extracted respectively are not in the same semantic space, and the relevant information of text and image is not noticed during the fusion process. Therefore these fused features cannot provide sufficient identification power to classify fake news. Dhruv et al. [21] proposed MVAE, which trains a decoder to reconstruct original text and low-level image features from the fused features. Although the ability of reconstruction means that the fused features can contain more information, the necessity of these auxiliary tasks in fake news detection remains unknown and computationally expensive. The MKEMN method proposed by Zhang et al. [7] combined aligned embeddings of text, image, and knowledge to learn multimodal representations for each post, applied in multimodal fake news detection. Wu et al. [6] proposed MCAN, which stacked multiple coattention layers to fuse multimodal features, allowing the learning of interdependencies between multiple modalities. Wei et al. [22] introduced CMC, where the network trained two single-modality networks through contrastive learning to learn cross-modal relevance. The aforementioned works mainly focus on feature-based fusion, thereby weakening the semantic relationships between modalities.

C. Multimodal Matching

Nowadays, semantics have been introduced into fake news detection, by some modal matching techniques, showing a novel fake news detection method. Zhou et al. [11] proposed SAFE, which fed the correlation between news textual and visual information into a classifier to detect fake news. However, the semantic gap in translating images to texts might have limited the ability to effectively utilize the cross-modal consistency. Xue et al. [12] introduced MCNN, which integrated textual semantic features, visual tampering features, and the similarity between textual and visual information for fake news detection. However, the aggregation process simply concatenated all features and did not consider the issue of cross-modal heterogeneity. Chen et al. [23] adopted a cross-modal alignment

method to train the encoder model, mapping text and visual data into a shared semantic space, and then utilize fused features for classification. Due to the limited number of training datasets and the adoption of suboptimal labeling methods, this may limit the effectiveness of the encoder, resulting in a semantic gap remaining between text and image features. This method uses cross-modal ambiguity scores to reweight multimodal features, but processing single-modal features through variational autoencoders with nonshared weights may affect the handling of ambiguity and thus model performance.

Based on the above related work, this article proposes a model that not only considers the consistency between multimodalities but also solves their heterogeneity, making it form a relationship of cooperation and competition. We also use attention-guided enhanced modal interaction and dynamic weight assignment based on similarity to create robust frameworks designed to enhance the detection of AI-generated fake news. Unlike traditional postfusion detection paradigms, this approach integrates text modality, visual modality (images), and embedded text within images (auxiliary modality) into a unified end-to-end framework. This framework encompasses four components: modality integration, feature extraction, feature interaction, and feature matching, for comprehensive news content analysis. By meticulously addressing the complex relationships between texts and images, this method improves detection accuracy and efficiency for fake news created by generative AI technologies.

III. METHODOLOGY

The objective of this study is to ascertain the authenticity of news sourced from generative AI-enabled social networks or other platforms, based on an analysis of its textual and visual components. Given a news article $A = (T, I, E)$ containing text (T), image (I), and image-embedded content (E). Our objective is to discern the authenticity of the news, classifying it as either fake ($y = 1$) or real ($y = 0$). This determination relies on evaluating the correlations represented by $S(T^I, I^T)$, which means the correlations between the textual and visual components, where $S(T, E)$ refers to the correlations between the original text and embedded content within the image, and $S(E^I, I^E)$ refers to the correlations between the image-embedded content and the image itself. The model under consideration is denoted as MFND-CMM, and its architectural representation is illustrated in Fig. 2.

A. Multimodal Feature Extraction

1) *Textual Feature Extraction Module*: For the extraction of textual features, this article employs a pretrained Bert [24] model comprising 12 encoder layers. This model is utilized to characterize news text, as represented by the following:

$$F_t^i = \text{Bert}(T^i), T = [T^0, T^1, \dots, T^{L-1}] \quad (1)$$

where T signifies the input sentence, L denotes the sentence length, and i signifies the i th word within the sentence. The term “Bert” encapsulates pretrained models designed for both Chinese and English languages.

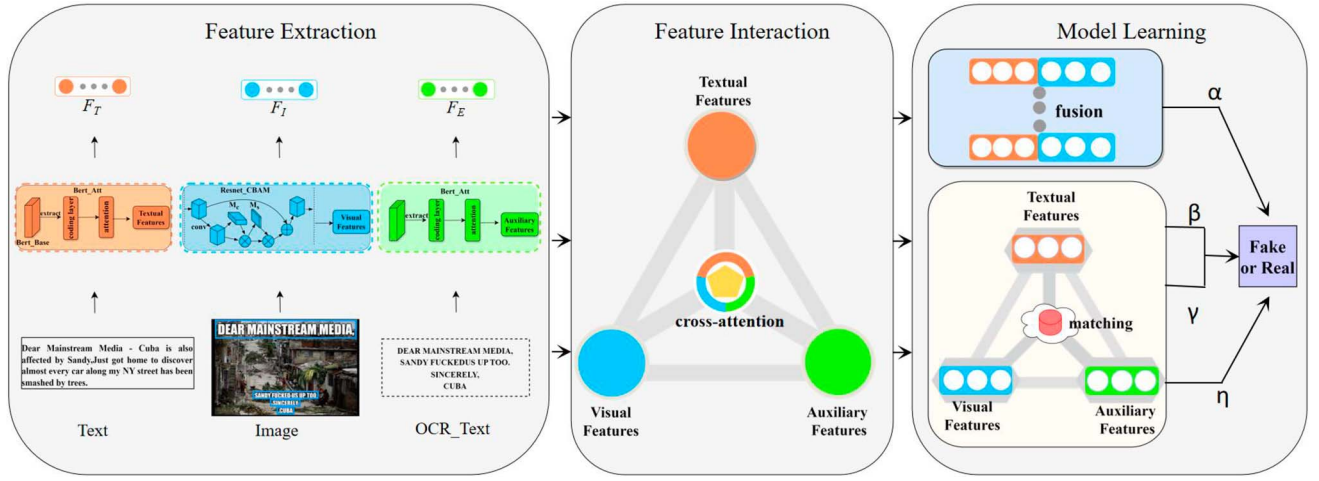


Fig. 2. MFND-CMM framework involves the utilization of three distinct encoders: Textual encoder for extracting textual features, visual encoder for extracting visual features, and auxiliary encoder for extracting image-embedded content, referred to as OCR_text, serving as auxiliary features. The extracted textual and visual features are amalgamated, and similarities among textual, visual, and auxiliary features are computed individually. Subsequently, various weights are assigned to facilitate a comprehensive and integrated decision-making process. The auxiliary modality is delineated as OCR_text in the diagram.

This article incorporates a self-attention mechanism [25] into the top layer of the Bert model. This augmentation enables the model to precisely capture the significance of each position within the input sequence, moving beyond a mere averaging of outputs from individual positions [26]. The calculation process is outlined as follows:

$$Q_s = (W_q F_t^i), K_s = (W_k F_t^i), V_s = (W_v F_t^i) \quad (2)$$

$$F_T = \text{SoftMax} \left(\frac{Q_s K_s^T}{\sqrt{d}} \right) V_s \quad (3)$$

where Q_s , K_s , and V_s are derived from textual features, W_q , W_k , and W_v represent weight matrices, d is the dimension of the input features, and F_T represents the textual features with attention. This feature extraction module is defined as Bert_Att.

2) *Visual Feature Extraction Module*: In this study, an enhanced ResNet-50 [27] model is chosen for visual feature extraction. When presented with an image I , the model processes it through ResNet to obtain an intermediate feature map F_i

$$F_i = \text{ResNet}(I). \quad (4)$$

The CBAM framework [28] integrates channel attention and spatial attention, enhancing ResNet's capability to extract more prominent features. This article incorporates CBAM into ResNet-50, denoting the modified model as ResNet_CBAM. Channel attention explicitly captures the interdependencies among feature channels, enabling the network to autonomously learn the significance of each channel. Consequently, the intermediate feature map F_i undergoes processing through channel attention, resulting in the refined feature map F_i'

$$M_c(F_i) = \sigma \{ \text{MLP}[\text{AvgPool}(F_i)] + \text{MLP}[\text{MaxPool}(F_i)] \} \quad (5)$$

$$F_i' = M_c(F_i) \otimes F_i \quad (6)$$

where σ represents the activation function, MLP is shared for both inputs, $M_c(F_i)$ is a 1-D channel attention map, and $\otimes$ means element-wise multiplication.

Spatial attention functions as an adaptive mechanism for selecting spatial regions, directing the model to focus more on key feature areas. Subsequently, the refined feature map F_i' is input into the spatial attention module

$$M_s(F_i') = \sigma \{ f^{7 \times 7} ([\text{AvgPool}(F_i'); \text{MaxPool}(F_i')]) \} \quad (7)$$

$$F_i'' = M_s(F_i') \otimes F_i' \quad (8)$$

$$F_I = F_i'' \oplus F_i \quad (9)$$

where σ represents the activation function, $f^{7 \times 7}$ denotes a convolution operation with a filter size of 7×7 , $M_s(F_i')$ is a 2-D spatial attention map, and F_I represents the visual features with attention.

3) *Image-Embedded Content Extraction Module*: Upon observation, it has been noted that numerous news article images encompass substantial embedded information. This embedded information commonly constitutes text and encapsulates the core theme of the news [29]. This study deems image-embedded information as an auxiliary modality. The open-source model, PP-OCRv3 [30], is employed for text extraction. Specifically, the pretrained PP-OCRv3 model is utilized in this article to extract the image-embedded content (E)

$$E = \text{OCR}(I) \quad (10)$$

where I refers to the image.

Likewise, Bert_Att is employed for feature extraction yielding vectorized representations and thereby producing the auxiliary feature F_E

$$F_E = \text{Bert_Att}(E^i), E = [E^0, E^1, \dots, E^{L-1}]. \quad (11)$$

B. Multimodal Feature Interaction

In addressing the interplay between textual and visual components in news articles, this section strives to establish alignment and interaction of information between the textual content and fine-grained images. To accomplish this, cross-attention

[31] is employed to facilitate the integration of cross-modal information between texts and images.

In the textual branch, the visual feature F_I undergoes mapping through three separate fully connected layers to serve as Q_{c1} , while the textual feature F_T is simultaneously mapped as both K_{c1} and V_{c1} . Following this, the cross-modal attention process is employed to generate the textual feature F_{T^I} , establishing correlation with the images

$$Q_{c1} = (W_q F_I), K_{c1} = (W_k F_T), V_{c1} = (W_v F_T) \quad (12)$$

$$F_{T^I} = \text{SoftMax} \left(\frac{Q_{c1} K_{c1}^T}{\sqrt{d}} \right) V_{c1}. \quad (13)$$

Similarly, within the image branch, the textual feature F_T is mapped through three separate fully connected layers to function as Q_{c2} , whereas the visual feature F_I is simultaneously mapped as both K_{c2} and V_{c2} . Following this, the cross-modal attention process is applied, resulting in the visual feature F_{I^T} that correlates with the texts

$$Q_{c2} = (W_q F_T), K_{c2} = (W_k F_I), V_{c2} = (W_v F_I) \quad (14)$$

$$F_{I^T} = \text{SoftMax} \left(\frac{Q_{c2} K_{c2}^T}{\sqrt{d}} \right) V_{c2}. \quad (15)$$

Via the cross-modal cross-attention process, the model has extracted crucial and pertinent information from the original textual feature F_T and visual feature F_I , denoted as F_{T^I} and F_{I^T} , respectively.

The cross-modal attention mechanism between the auxiliary modality and image is analogous to the cross-modal attention between text and image, so it will not be reiterated here. Consequently, the model has extracted significant relevant information denoted as F_{E^I} and F_{I^E} from the original cross-modal feature F_E and visual feature F_I , respectively.

C. Multimodal Feature Fusion

Up to this point, textual and visual features have been acquired. Subsequently, the final dimension is obtained by concatenating their last dimensions without explicitly considering their interrelationship, aiming to accurately predict the likelihood that the news is false. This can be mathematically defined as

$$C = \text{SoftMax}(W_C(F_{T^I} \oplus F_{I^T}) + B_C) \quad (16)$$

where $\oplus$ is the symbol of splicing operation, C represents the probability of class prediction, W_C and B_C are weight parameters and offset items respectively.

To make the calculated possibility of falsification of a news article close to its true label y , a loss function based on cross-entropy is defined as

$$\mathcal{L}_c(\theta_{T^I}, \theta_{I^T}, \theta_c) = -\mathbb{E}_{(a,y) \sim (A,Y)} [y \cdot \log C + (1-y) \cdot \log(1-C)] \quad (17)$$

where Y is the set of news true and fake label categories, A is the news set, and y refers to the actual label of news a . θ_{T^I} , θ_{I^T} , and θ_c respectively represent the corresponding parameter matrices.

D. Multimodal Feature Matching

The correlation between the textual and visual elements of news articles can serve as a means to assess the authenticity of news stories. Creators of fake news may intentionally select unrelated images as a basis for deceptive claims, aiming to capture readers' attention.

This article assesses the concordance between the two modalities utilizing a slightly modified cosine similarity for simplicity in calculation and precise measurement. Initially, examine the coherence of the textual feature F_{T^I} and visual feature F_{I^T} , defined mathematically as follows:

$$S(F_{T^I}, F_{I^T}) = \frac{F_{T^I} \cdot F_{I^T} + \|F_{T^I}\| \|F_{I^T}\|}{2 \|F_{T^I}\| \|F_{I^T}\|}. \quad (18)$$

Similarly, when examining the coherence between the auxiliary feature F_{E^I} and the visual feature F_{I^E} , their mathematical definition aligns with the calculation of consistency between textual and visual features. Consequently, their similarity can be expressed as $S(F_{E^I}, F_{I^E})$. In evaluating the coherence between the original textual feature F_T and the original auxiliary feature F_E , as they belong to the same modality, cross-modal attention interaction is unnecessary. This is mathematically defined as $S(F_T, F_E)$.

Then, define the following three cross-entropy loss functions to represent that from the perspective of similarity:

$$\mathcal{L}_s(\theta_{T^I}, \theta_{I^T}) = -\mathbb{E}_{(a,y) \sim (A,Y)} \{y \cdot \log[1 - S(F_{T^I}, F_{I^T})] + (1-y) \cdot \log S(F_{T^I}, F_{I^T})\} \quad (19)$$

The mathematical definitions of $\mathcal{L}_s(\theta_{E^I}, \theta_{I^E})$ and $\mathcal{L}_s(\theta_T, \theta_E)$ are the same as $\mathcal{L}_s(\theta_{T^I}, \theta_{I^T})$ and will not be repeated here.

E. Multimodal Model Learning

Ultimately, the correct identification of fake news is achieved through the fusion and assessment of similarities among multimodal features of news. Here, α , β , γ , and η represent the loss function weights for the four branch networks. To incorporate the significance of all four modalities, the final loss function formulations are as follows:

$$\begin{aligned} \mathcal{L}(\theta_T, \theta_E, \theta_c, \theta_{T^I}, \theta_{I^T}, \theta_{E^I}, \theta_{I^E}) \\ = \alpha \mathcal{L}_c(\theta_{T^I}, \theta_{I^T}, \theta_c) + \beta \mathcal{L}_s(\theta_{T^I}, \theta_{I^T}) \\ + \gamma \mathcal{L}_s(\theta_T, \theta_E) + \eta \mathcal{L}_s(\theta_{E^I}, \theta_{I^E}) \end{aligned} \quad (20)$$

$$\begin{aligned} (\hat{\theta}_T, \hat{\theta}_E, \hat{\theta}_c, \hat{\theta}_{T^I}, \hat{\theta}_{I^T}, \hat{\theta}_{E^I}, \hat{\theta}_{I^E}) \\ = \text{argmin} \mathcal{L}(\theta_T, \theta_E, \theta_c, \theta_{T^I}, \theta_{I^T}, \theta_{E^I}, \theta_{I^E}). \end{aligned} \quad (21)$$

IV. EXPERIMENTS

A. Datasets

To comprehensively assess the performance of MFND-CMM, experiments are conducted on two publicly available datasets. Initial statistics reveal that over 70% of the images in both datasets incorporate embedded text. The dataset statistics are presented in Table I. The following provides a description of the datasets:

TABLE I
DATASET STATISTIC

Dataset	Label	Number	All
Weibo	fake	4749	9528
	real	4779	
Twitter	fake	7021	12995
	real	5974	

1) The Weibo dataset [32] is sourced from the Weibo social platform. This Weibo dataset of true information is collected from authoritative Chinese sources, and fake information is obtained through the official Weibo rumor suppression system. The dataset contains 4749 pieces of fake news data and 4779 pieces of real news data, along with 9528 news images. Each data instance includes both a news text and a corresponding news image. 2) The Twitter dataset [33] is sourced from the Twitter social platform. The content of each tweet has a brief text message and extra images or videos. There are around 6000 rumors and 5000 nonrumor tweets in the development set from 11 rumor-related events. It includes 7021 pieces of fake news data, and 5974 pieces of real news data, and 517 news images.

They are split into 7:1:2 ratios for training, validation, and test sets.

B. Experimental Settings

The dimensionality of textual features obtained from Bert_Att is 768. Images are first resized to $224 \times 224 \times 3$ and then fed to ResNet_CBAM. The dimensionality of visual features obtained from ResNet_CBAM is 2048. The ultimately obtained cross-modal features are all of dimensions 32×32 .

In experiments, Adam is selected as the optimizer. The batch size is 32, and the epoch is 100. Weibo's learning rate is set to 0.001, whereas Twitter's learning rate is set to 0.0001, and 0.5 is the dropout rate. After comparing several groups of experiments and assigning weights according to the importance of branches and manual experience, the weights of the four network branches in the model are finally determined as $\alpha = 0.6$, $\beta = 0.2$, $\gamma = 0.1$, and $\eta = 0.1$.

In the setup of metrics for experimental results, precision refers to the proportion of samples that are actually positive among those predicted as positive by the model. It measures the accuracy of the model's predictions. Recall refers to the proportion of samples that are actually positive and are correctly predicted as positive by the model. It measures the model's ability to identify positive samples. The F1 Score is the harmonic mean of precision and recall, taking into account the importance of both precision and recall. It is a comprehensive metric that combines the two.

C. Baselines

- 1) *att-RNN* [5]: The att-RNN model incorporates unidirectional LSTM and VGG-19 for the extraction of image features, context features, and textual features. Subsequently, it employs an RNN with attention mechanisms to discern and identify fake news.

- 2) *EANN* [8]: EANN employs Text-CNN for textual feature extraction and VGG for visual feature extraction. It then obtains news features by concatenating textual and visual information, utilizing an event discriminator to detect fake news.
- 3) *MVAE* [21]: MVAE utilizes bidirectional LSTM for textual feature extraction and VGG for visual feature extraction. These features are then fused into an autoencoder to reconstruct the correlations between different modes of feature learning.
- 4) *Spotfake* [20]: Spotfake utilizes Bert for textual feature extraction and VGG-19 for visual feature extraction. These features are subsequently concatenated to facilitate the detection of fake news.
- 5) *MKEMN* [7]: MKEMN considers texts, images, and retrieved knowledge embeddings as stacked channels and performs fusion through a convolutional operation.
- 6) *BDANN* [34]: BDANN uses BERT to extract text features and VGG-19 to extract image features. It also applies domain classifiers to eliminate dependencies between features, which is ultimately used for news detection.
- 7) *SAFE* [11]: SAFE converts the image into text and extracts textual features using Text-CNN. Finally, cosine similarity is employed to detect the similarity between textual representations.
- 8) *MCNN* [12]: MCNN integrates textual semantic features, visual tampering features, and the similarity of textual and visual information computed through cosine similarity for use in fake news detection.
- 9) *MCAN* [6]: MCAN stacks multiple co-attention layers to fuse the multimodal features.
- 10) *CAFE* [23]: CAFE quantifies cross-modal ambiguity to adaptively aggregate unimodal features and cross-modal correlations.
- 11) *LIIMR* [35]: LIIMR identifies and suppresses information from weaker modalities and extracts relevant information from the strong modality on a per-sample basis.
- 12) *MRHFR* [36]: MRHFR integrates multimodal news features by simulating three reading habits of humans, and utilizes a consistency-constrained reasoning layer to detect inconsistencies between different modalities, thereby achieving multimodal fake news detection.

D. Performance Comparison

Table II juxtaposes the performance of MFND-CMM against baseline models. Notably, MFND-CMM attains the pinnacle of accuracy on both real-world datasets, registering impressive values of 90.3% and 93.2%, respectively. Particularly noteworthy is the F1 score in this study, which stands out as the highest among all metrics. Given that F1 embodies a delicate equilibrium between precision and recall, the overall evaluation underscores that this article attains the most favorable metrics.

From the perspective of feature extraction, Spotfake and MKEMN models adopt pretrained feature extractors without considering the relative importance of each feature. Experimental results on the Weibo dataset reveal that MFND-CMM

TABLE II
COMPARISON OF EXPERIMENTAL RESULTS FOR METHODS

Dataset	Method	Accuracy	Real Information			Fake Information		
			Precision	Recall	F1	Precision	Recall	F1
Weibo	att-RNN [5]	0.788	0.738	0.890	0.807	0.862	0.686	0.764
	EANN [8]	0.816	0.810	0.810	0.810	0.820	0.820	0.820
	MVAE [21]	0.824	0.802	0.875	0.837	0.854	0.769	0.809
	Spotfake [20]	0.892	0.847	0.656	0.739	0.902	0.964	0.932
	MKEMN [7]	0.824	0.723	0.819	0.798	0.823	0.799	0.812
	BDANN [34]	0.842	0.850	0.820	0.830	0.830	0.870	0.850
	SAFE [11]	0.816	0.816	0.818	0.817	0.818	0.815	0.817
	MCNN [12]	0.823	0.787	0.848	0.816	0.858	0.801	0.828
	MCAN [6]	0.899	0.884	0.909	0.897	0.913	0.889	0.901
	CAFE [23]	0.840	0.825	0.851	0.837	0.855	0.830	0.842
	LIIMR [35]	0.900	0.882	0.823	0.847	0.908	0.941	0.925
	MRHFR [36]	0.907	0.939	0.869	0.903	0.879	0.931	0.904
	MFND-CMM	0.903	0.885	0.912	0.898	0.914	0.923	0.918
Twitter	att-RNN [5]	0.682	0.603	0.770	0.676	0.780	0.615	0.689
	EANN [8]	0.719	0.771	0.870	0.817	0.642	0.474	0.545
	MVAE [21]	0.745	0.689	0.777	0.730	0.801	0.719	0.758
	Spotfake [20]	0.777	0.832	0.606	0.701	0.751	0.900	0.820
	MKEMN [7]	0.714	0.634	0.814	0.831	0.814	0.756	0.708
	BDANN [34]	0.830	0.830	0.930	0.880	0.810	0.630	0.710
	SAFE [11]	0.762	0.695	0.811	0.748	0.831	0.724	0.774
	MCNN [12]	0.784	0.790	0.787	0.788	0.778	0.781	0.779
	MCAN [6]	0.809	0.732	0.871	0.795	0.889	0.765	0.822
	CAFE [23]	0.806	0.805	0.813	0.809	0.807	0.799	0.803
	LIIMR [35]	0.831	0.836	0.832	0.830	0.825	0.830	0.827
	MRHFR [36]	0.921	0.976	0.828	0.896	0.876	0.981	0.926
	MFND-CMM	0.932	0.922	0.941	0.931	0.922	0.921	0.922

Note: The best indexes are in bold and the second best ones are in red.

outperforms both Spotfake and MKEMN, achieving an accuracy that is 2.1% higher than Spotfake and 7.9% higher than MKEMN. This substantiates the assertion that the MFND-CMM model discerns the relative importance of textual and image features through attention levels, thereby enhancing the efficacy of feature representation and, consequently, improving the effectiveness of fake news detection. From the perspective of feature fusion, the MFND-CMM model exhibited a notable enhancement in accuracy on both the Weibo dataset (7.9% increase) and the Twitter dataset (18.7% increase) when compared with the MVAE model without a fusion mechanism. This underscores the efficacy of the feature fusion mechanism within the MFND-CMM model, aiding both modalities in uncovering more crucial information and effectively elevating the expressiveness of modal features.

From the perspective of feature matching, the absence of cross-modal interaction between the SAFE model and MCNN model results in a challenge of modal heterogeneity. Experimental outcomes on the Weibo dataset reveal that MFND-CMM surpasses both SAFE and MCNN, showcasing an accuracy rate that is 8.7% higher than SAFE and 8% higher than MCNN. This substantiates that the MFND-CMM model maximizes the utilization of multimodal interaction data, enhancing the association features between text and image, and effectively addressing the issue of modal heterogeneity. Furthermore, MCAN does not pay enough attention to the consistency between features, and CAFE has limitations on single-modal features, which will affect the robustness of the model.

To further prove the validity and correctness of the model proposed in the text, the training process during the

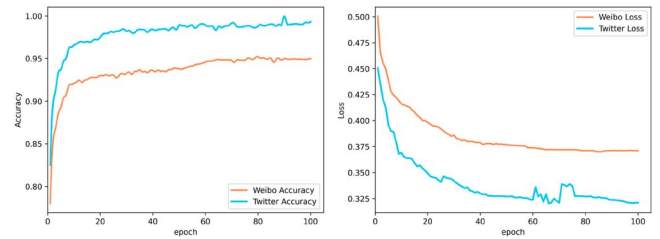


Fig. 3. Visualization results of the experimental training process on two real-world datasets.

experiment is retained and visualized in this article, as shown in Fig. 3.

E. Architecture Ablation Analysis

The article conducts a series of comparative experiments employing the original Bert and ResNet to elucidate the efficacy of the proposed feature extractor. Furthermore, component ablation experiments are executed to underscore the effectiveness of each constituent of MFND-CMM. The resultant experimental findings encompass evaluation metrics for the fake Information class, encompassing accuracy, precision, recall, and F1 score.

1) *Ablation Experiments of Feature Extractor:* The pertinent experimental outcomes are depicted in Figs. 4 and 5. Notably, the original pre-trained feature extractor yields suboptimal results. Conversely, the enhanced model exhibits the most favorable experimental outcomes, showcasing the efficacy of the proposed feature extractor. The discernible salient features extracted aptly capture the distinctive characteristics of news,

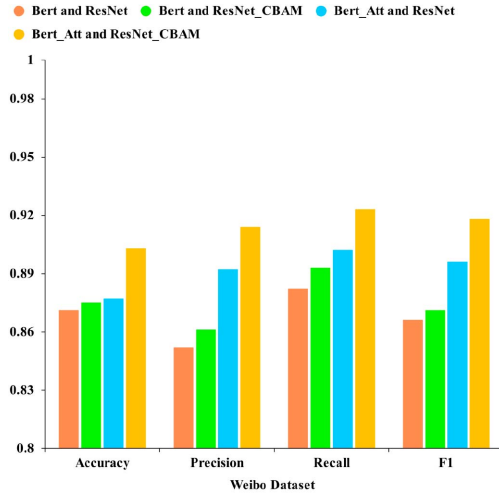


Fig. 4. Architecture ablation analysis of feature extractor.

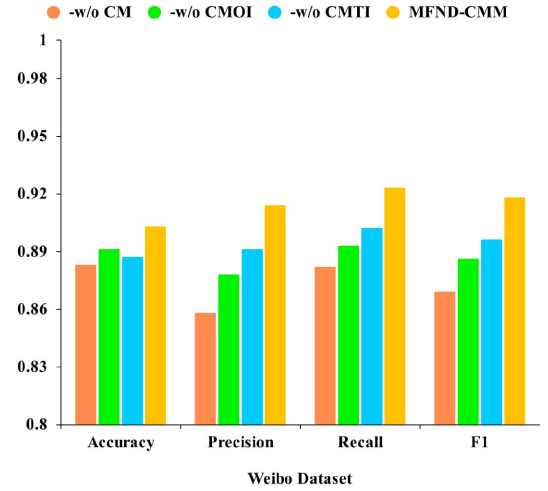


Fig. 6. Architecture ablation analysis of feature interaction.

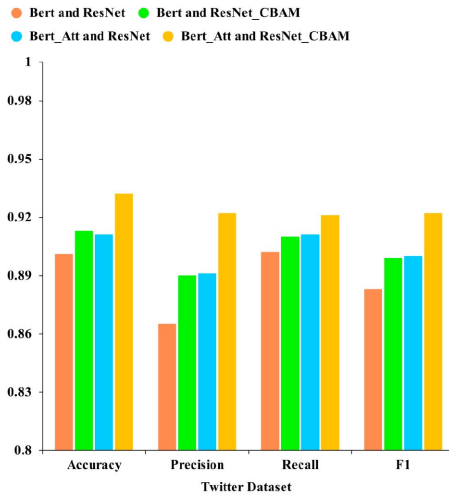


Fig. 5. Architecture ablation analysis of feature extractor.

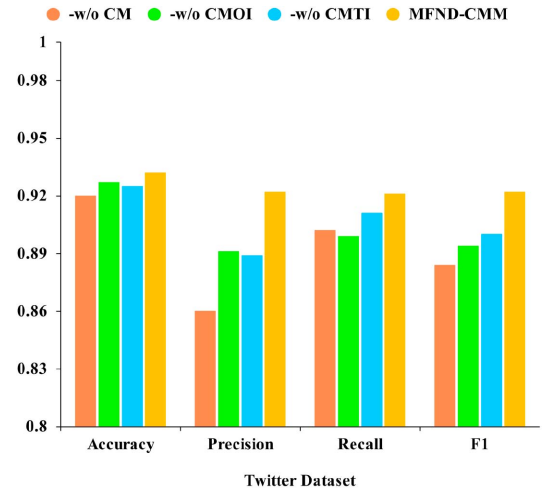


Fig. 7. Architecture ablation analysis of feature interaction.

further validating the effectiveness of the feature extractor delineated in this article.

2) Ablation Experiments of Feature Interaction:

- 1) -w/o CM: Experiments with no cross-modal.
- 2) -w/o CMOI: Remove cross-modal updates between OCR_text and image. Conduct experiments with cross-modal interactions between text and image.
- 3) -w/o CMTI: Remove cross-modal updates between text and image. Conduct experiments with cross-modal interactions between OCR_text and image.
- 4) MFND-CMM: Experiments with cross-modal.

The pertinent experimental outcomes are illustrated in Figs. 6 and 7. The cross-attention mechanism facilitates dynamic integration between modalities, wherein visual updates retain pivotal textual information, and textual updates preserve essential visual details. This concurrent process effectively mitigates heterogeneity across modalities. Consequently, it can be inferred that the feature interaction module constructed in this article proves to be effective.

3) Ablation Experiments of Individual Components of MFND-CMM:

- 1) -w/o $\alpha\beta\gamma$: Experiments with image and image-embedded content matching.
- 2) -w/o $\alpha\beta\eta$: Experiments with text and image-embedded content matching.
- 3) -w/o $\alpha\gamma\eta$: Experiments with text-image matching.
- 4) -w/o $\beta\gamma\eta$: Experiments with text-image fusion.
- 5) -w/o $\gamma\eta$: Text-image fusion, text-image matching, experiments with weights of 0.6 and 0.4 respectively.
- 6) -w/o η : Text-image fusion, text, and image content matching, image and image-embedded content matching, experiments with weights of 0.6, 0.2, and 0.2, respectively.

Figs. 8 and 9 reveal that text-image fusion exerts the most substantial influence on fake news detection. The semantic consistency achieved through similarity-based dynamic weight assignment emerges as a pivotal factor in enhancing the model's performance. It is evident that mere text-image fusion and

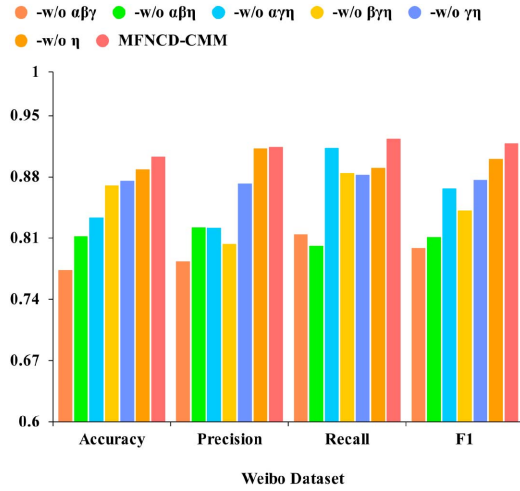


Fig. 8. Architecture ablation analysis of MFND-CMM.

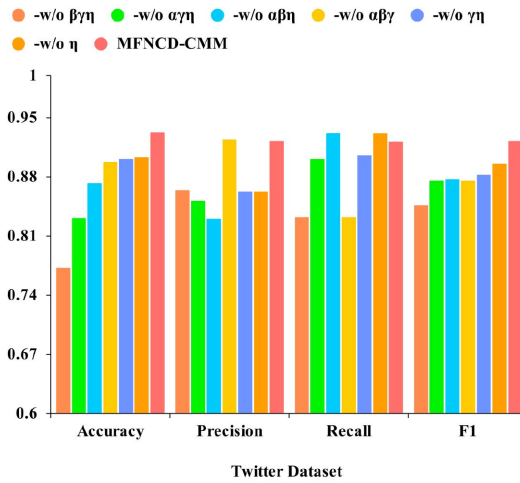


Fig. 9. Architecture ablation analysis of MFND-CMM.

text–image matching fall short of achieving the desired impact. However, with the incorporation of image-embedded content, a notable improvement in the indicator is observed. This underscores the effectiveness of the text–image consistency detection method and underscores the significance of modeling image-embedded content.

F. Case Study

In our case study, we seek to substantiate the following inquiries: Do instances of fake news exist in the real world where the texts and images exhibit inconsistency? Can our model adeptly discern and accurately identify this incongruence, thus flagging the news as falsified? In pursuit of these objectives, we scrutinized news articles from two datasets, aligning their ground truth labels with the predicted labels from our model. Our findings indicate that our model adeptly identifies the consistent relationship between texts and images, serving as a pivotal criterion for discerning authenticity. Fig. 10 shows cases several examples, revealing a discernible correspondence



Fig. 10. Cases of fabricated news and real news in real-world datasets successfully identified by the MFND-CMM model. (a) Fake news. (b) Real news.

between the textual and visual components of genuine news, while the text of fake news struggles to find support from unaltered images.

Future research could consider improvements in the following areas:

- 1) *Cross-Modal Information Integration*: Current multi-modal research primarily focuses on the integration of news text and images. Future directions could be dedicated to exploring more efficient ways to merge various modalities such as text, images, and videos.
- 2) *System Upgrades*: In light of the rapid dissemination and dynamic nature of fake news, future studies should focus on developing a fake news detection system capable of real-time monitoring and dynamic response.

V. CONCLUSION

This article suggests a novel multimodal fake news detection framework through cross-modal matching, which holds the potential for application in future social networks. The proposed framework comprises a feature extraction module, a feature interaction module, a fusion module, and three matching modules. The feature extraction module captures essential features from news articles. In the fusion module, textual, and visual features are merged. The matching modules evaluate intermodality correlations using pairwise-similarity scores. Simultaneously, the feature interaction module utilizes the cross-attention mechanism to introduce dynamic integration across modalities. To strike a balance between consistency and heterogeneity, attention-guided modality-enhanced interaction and

similarity-based dynamic weight assignment are used, contributing to the creation of a robust framework. Experiments show that our method outperforms existing baseline methods. In the future work, a robust framework aimed at improving the detection of AI-generated fake news is still open.

REFERENCES

- [1] L. Tan, G. Wang, F. Jia, and X. Lian, "Research status of deep learning methods for rumor detection," *Multimedia Tools Appl.*, vol. 82, no. 2, pp. 2941–2982, 2023.
- [2] C. Castillo, M. Mendoza, and B. Poblete, "Information credibility on Twitter," in *Proc. Web Conf.*, 2011.
- [3] J. Ma et al., "Detecting rumors from microblogs with recurrent neural networks," in *Proc. 25th Int. Joint Conf. Artif. Intell. (IJCAI)*, 2016, pp. 3818–3824.
- [4] V. Qazvinian, E. Rosengren, D. R. Radev, and Q. Mei, "Rumor has IT: Identifying misinformation in microblogs," in *Proc. Empirical Methods Nat. Lang. Process.*, 2011.
- [5] Z. Jin, J. Cao, H. Guo, Y. Zhang, and J. Luo, "Multimodal fusion with recurrent neural networks for rumor detection on microblogs," in *Proc. ACM*, 2017.
- [6] Y. Wu, P. Zhan, Y. Zhang, L. Wang, and Z. Xu, "Multimodal fusion with co-attention networks for fake news detection," in *Proc. Findings ACL*, 2021.
- [7] H. Zhang, Q. Fang, S. Qian, and C. Xu, "Multi-modal knowledge-aware event memory network for social media rumor detection," in *Proc. 27th ACM Int. Conf.*, 2019.
- [8] Y. Wang et al., "EANN: Event adversarial neural networks for multi-modal fake news detection," in *Proc. ACM SIGKDD*, 2018, pp. 849–857.
- [9] Y. J. Lu and C. T. Li, "GCAN: Graph-aware co-attention networks for explainable fake news detection on social media," 2020, *arXiv:2004.11648*.
- [10] C. Zhang, Z. Yang, X. He, and L. Deng, "Multimodal intelligence: Representation learning, information fusion, and applications," *IEEE J. Sel. Topics Signal Process.*, vol. 14, no. 3, pp. 478–493, Mar. 2020.
- [11] X. Zhou, J. Wu, and R. Zafarani, "Safe: similarity-aware multi-modal fake news detection," 2020, *arXiv:200304981*.
- [12] J. Xue, Y. Wang, Y. Tian, Y. Li, L. Shi, and L. Wei, "Detecting fake news by exploring the consistency of multimodal data," *Inf. Process. Manage.*, vol. 58, no. 5, 2021, Art. no. 102610.
- [13] J. Chen, C. Jia, H. Zheng, R. Chen, and C. Fu, "Is multi-modal necessarily better? robustness evaluation of multi-modal fake news detection," *IEEE Trans. Netw. Sci. Eng.*, vol. 10, no. 6, pp. 3144–3158, Nov./Dec. 2023.
- [14] F. Qian, C. Gong, K. Sharma, and Y. Liu, "Neural user response generator: Fake news detection with collective user intelligence," in *Proc. 27th Int. Joint Conf. Artif. Intell. (IJCAI)*. International Joint Conferences on Artificial Intelligence Organization, 2018, pp. 3834–3840, doi: 10.24963/ijcai.2018/533.
- [15] B. Bhattarai, O. C. Granmo, and L. Jiao, "Explainable Tsetlin machine framework for fake news detection with credibility score assessment," 2021, *arXiv:2105.09114*.
- [16] A. Gupta, H. Lamba, P. Kumaraguru, and A. Joshi, "Faking sandy: characterizing and identifying fake images on Twitter during hurricane sandy," in *Proc. 22nd Int. Conf. World Wide Web*, 2013, pp. 729–736.
- [17] P. Qi, J. Cao, T. Yang, J. Guo, and J. Li, "Exploiting multi-domain visual information for fake news detection," in *Proc. IEEE Int. Conf. Data Mining (ICDM)*, Piscataway, NJ, USA: IEEE Press, 2019, pp. 518–527.
- [18] J. Cao, P. Qi, Q. Sheng, T. Yang, and J. Li, "Exploring the role of visual content in fake news detection," in *Disinformation, Misinformation, and Fake News in Social Media: Emerging Research Challenges and Opportunities*, 2020, pp. 141–161.
- [19] H. Zhang, Q. Fang, S. Qian, and C. Xu, "Multi-modal knowledge-aware event memory network for social media rumor detection," in *Proc. ACM Int. Conf. Multimedia*, 2019, pp. 1942–1951.
- [20] S. Singhal, R. R. Shah, T. Chakraborty, P. Kumaraguru, and S. Satoh, "Spotfake: A multi-modal framework for fake news detection," in *Proc. IEEE 5th Int. Conf. Multimedia Big Data (BigMM)*, Piscataway, NJ, USA: IEEE Press, 2019, pp. 39–47.
- [21] D. Khattar, J. S. Goud, M. Gupta, and V. Varma, "MVAE: Multimodal variational autoencoder for fake news detection," in *Proc. World Wide Web Conf.*, 2019, pp. 2915–2921.
- [22] Z. Wei, H. Pan, L. Qiao, X. Niu, P. Dong, and D. Li, "Cross-modal knowledge distillation in multi-modal fake news detection," in *Proc. IEEE Int. Conf. Acoust., Speech Signal Process. (ICASSP)*, Piscataway, NJ, USA: IEEE Press, 2022, pp. 4733–4737.
- [23] Y. Chen et al., "Cross-modal ambiguity learning for multimodal fake news detection," in *Proc. ACM Web Conf.*, 2022, pp. 2897–2905.
- [24] J. Devlin, M.-W. Chang, K. Lee, and K. Toutanova, "BERT: Pre-training of deep bidirectional transformers for language understanding," 2018, *arXiv:1810.04805*.
- [25] A. Vaswani et al., "Attention is all you need," *Adv. Neural Inf. Process. Syst.*, 2017, p. 30.
- [26] Z. Lin et al., "A structured self-attentive sentence embedding," 2017, *arXiv:1703.03130*.
- [27] K. He, X. Zhang, S. Ren, and J. Sun, "Deep residual learning for image recognition," in *Proc. IEEE Conf. Comput. Vis. Pattern Recognit.*, 2016, pp. 770–778.
- [28] S. Woo, J. Park, J.-Y. Lee, and I. S. Kweon, "CBAM: Convolutional block attention module," in *Proc. ECCV*, 2018, pp. 3–19.
- [29] P. Qi et al., "Improving fake news detection by using an entity-enhanced framework to fuse diverse multimodal clues," in *Proc. 29th ACM Int. Conf. Multimedia*, 2021, pp. 1212–1220.
- [30] C. Li et al., "PP-OCRv3: More attempts for the improvement of ultra lightweight OCR system," 2022, *arXiv:2206.03001*.
- [31] M. Gheini, X. Ren, and J. May, "Cross-attention is all you need: Adapting pretrained transformers for machine translation," 2021, *arXiv:2104.08771*.
- [32] Z. Jin, J. Cao, Y. Zhang, J. Zhou, and Q. Tian, "Novel visual and statistical image features for microblogs news verification," *IEEE Trans. multimedia*, vol. 19, no. 3, pp. 598–608, Mar. 2017.
- [33] Z. Jin, J. Cao, Y. Zhang, and Y. Zhang, "MCG-ICT at mediaeval 2015: Verifying multimedia use with a two-level classification model," in *Proc. MediaEval*, 2015.
- [34] T. Zhang et al., "BDANN: Bert-based domain adaptation neural network for multi-modal fake news detection," in *Proc. Int. Joint Conf. Neural Netw. (IJCNN)*, Piscataway, NJ, USA: IEEE Press, 2020, pp. 1–8.
- [35] S. Singhal, T. Pandey, S. Mrig, R. R. Shah, and P. Kumaraguru, "Leveraging intra and inter modality relationship for multimodal fake news detection," in *Proc. Companion Proc. Web Conf.*, 2022, pp. 726–734.
- [36] L. Wu, P. Liu, and Y. Zhang, "See how you read? multi-reading habits fusion reasoning for multi-modal fake news detection," in *Proc. AAAI Conf. Artif. Intell.*, vol. 37, no. 11, 2023, pp. 13736–13744.