

A Review on Media Forensics using Image Processing Techniques and Deep learning Tools

Dr. Chanda V Reddy¹, Anusha H², Dhanush N³, Madhushree T P⁴, Nischay P⁵

Professor, Head of Department, Department of Telecommunication, KSIT, Bengaluru, India¹

UG Student, Department of telecommunication, KSIT, Bengaluru, India²⁻⁵

Abstract: In the present generation, social media is a big advantage for an individual to grow. On the other hand, we can't neglect the fact that it's a huge platform for negativity too. With the rapid progress of recent years, techniques that generate and manipulate multimedia content can now provide a very advanced level of realism. The boundary between real and synthetic media has become very thin. On the one hand, this opens the door to a series of exciting applications in different fields such as creative arts, advertising, film production, video games. On the other hand, it poses enormous security threats. Software packages freely available on the web allow any individual, without special skills, to create very realistic fake images and videos. These techniques can be used to manipulate public opinion regarding anything and create chaos. In this paper, we would like to overview few major facts and figures regarding exceeding image forgery techniques that exists and propose a better way on how to detect these forgeries and fakes.

Keywords: GAN, ELA, deep learning, convolutional neural networks, Fake colorized image detection.

I. INTRODUCTION

Fake multimedia has become a huge problem in recent years since the development of many image processing and deep learning tools. With these tools, creating realistic so-called deep fakes and fake media is very much easy. These fake images are used to manipulate public opinion by fake news campaigns and also can be used for malicious purposes, like creating fake porn videos to blackmail people. Due to this, people are losing faith in journalism. Some fakes are easy to identify since they are made for fun and includes famous personalities. However, verifying digital integrity becomes much more difficult if the video portrays a less known person and only the manipulated version is publicly available. This scenario takes place, for example, if the attacker films a new video on his own, with a collaborative actor whose face is eventually replaced by the face of the targeted person. Governmental bodies, enforcement agencies, the news industry, and also the man in the street are becoming acutely aware of the potential menace carried by such a technology. The scientific community is asked to develop reliable tools for automatically detecting fake multimedia. Image manipulation has been carried out since photography was born², and powerful image/video editing tools, such as Photoshop® After Effects Pro®, or the open-source software GIMP, have been around for a long time. Using such conventional tools images can be easily modified, obtaining realistic results that can fool even a careful observer. In the literature survey below, many method have been proposed to detect these fakes and forgeries. Every image and video is characterized by numerous features. which depend on the different phases of its digital history: from the very same acquisition process, to the internal camera processing (e.g., demosaicing, compression), to all external processing and editing operations. Digital manipulations tend to modify such features, leaving a trail of clues which, although invisible to the eye, can be exploited by pixel-level analysis tools. Instead, semantic integrity is violated when the media asset under analysis conveys information which is not coherent with the context or with evidence coming from correlated sources. For example, when objects are copy-pasted from images available on the web, several near-identical copies can be detected, suggesting a possible manipulation. Moreover, by identifying the connections among the various versions of the same asset, it is possible to build its manipulation history. Despite the continuous research efforts and the numerous forensic tools developed in the past, the advent of deep learning is changing the rules of the game and asking multimedia forensics for new and timely solutions. This phenomenon is also causing a strong acceleration in multimedia forensics research, which often relies itself on deep learning.

II. LITERATURE SURVEY

[1] Quantization is the critical step in lossy compression which maps the DCT coefficients in an irreversible way under the quantization constraint set (QCS) theorem. In this paper, they first derive that a doubly compressed image no longer follows the quantisation constraint set (QCS) theorem and then propose a novel quantization noise model to characterize single and doubly compressed images. In order to detect double compression forgery, they further propose to approximate the uncompressed ground truth image using image restoration techniques. In this paper, they conduct a series of experiments to demonstrate the validity of the proposed quantization noise model and also the effectiveness of

the forgery detection method with the proposed image restoration techniques double compression forgeries. They have shown that the proposed quantization noise model indeed characterizes the change on compression characteristics before and after recompression and have justified the effectiveness of the proposed model with promising experimental results. The contributions of this work include a theoretical and content-independent model is proposed to detect double compression forgeries and the image restoration techniques are adopted to resolve the practical forgery detection problem. Using this restoration perspective, many other image acquisition properties could also be involved in this framework. In addition, the proposed approach can successfully locate the forged region as small as 8x8 block, either with aligned or misaligned block boundary cases. With these advantages, we believe the proposed model could also be combined with other existing forensic features to solve more complex forgery problems on compressed images.

To verify the robustness of the proposed quantization noise model, here they first assume that the un-quantized ground truth images are available during the forgery detection. They also compare with two existing approaches i.e., through measuring inconsistencies of blocking artifacts and exposing digital forgeries from JPEG, where the 1st method relies on estimation of primary quantization table and the 2nd method relies on exhaustively recompressing the test image using every possible primary quality factor, and both are able to locate forged regions. In order to have a fair comparison, here they assume the primary quantization table and quality factor are known. In this experiment, each image size is 1024x1024. They first crop a sub image with size 480x480 and compress this sub image with JPEG quality factor 1 QF. Next, they copy this compressed sub image back into the original raw image and then compress the spliced image with JPEG quality factor 2 QF. We test 500 images for each quality setting 1 QF and 2 QF and derive a ROC curve for each test image. Fig. 7 shows the averaged ROC curve when 1 QF equals to 50. The proposed quantization noise model achieves high performance in detecting the copy-paste-recompression forgery. Using their proposed framework, they approximate a reliable ground truth image and apply it to forgery detection. This paper shows a compressed image before image restoration, and also shows the difference between a compressed image with quality factor 80 and its restored result. The artifacts around block boundaries are now eliminated and the image details are also enhanced. This paper also shows the detection result, where the forged region is first compressed with quality factor 50 and then recompressed with quality factor 80. They also indicate the posterior map, where the forged region is clearly identified.

[2] Copy-move forgery is one of the most commonly used manipulations for tampering digital images. Key point-based detection methods have been reported to be very effective in revealing copy-move evidence due to their robustness against various attacks, such as large-scale geometric transformations. However, these methods fail to handle the cases when copy-move forgeries only involve small or smooth regions, where the number of key points is very limited. To tackle this challenge, in this paper they propose a fast and effective copy-move forgery detection algorithm through hierarchical feature point matching. They first show that it is possible to generate a sufficient number of key points that exist even in small or smooth regions by lowering the contrast threshold and rescaling the input image. We then develop a novel hierarchical matching strategy to solve the key point matching problems over a massive number of key points. To reduce the false alarm rate and accurately localize the tampered regions, they further propose a novel iterative localization technique by exploiting the robustness properties (including the dominant orientation and the scale information) and the color information of each key point. Extensive experimental results are provided to demonstrate the superior performance of their proposed scheme in terms of both efficiency and accuracy.

[3] In Recent years, Representation learning as one of the information extraction and data mapping methods in machine learning systems have received huge attention. Artificial deep neural networks are considered as one of the basic structures capable of representation learning. However, a large number of standard representations learning methods are supervised and requires a lot of labelled data. In this paper, they introduce an unsupervised representation learning by designing and implementing deep neural networks (DNNs) in combination with Generative Adversarial Networks (GANs). The main idea behind the proposed method, which causes the superiority of this method over others is representation learning via the generative models and encoder networks altogether. In this research, encoders are utilized in addition to the generative models to help the more features to be extracted. It is shown that the proposed method not only help feature extraction but accelerate and improve the performance of the learning in GANs which lead to better feature extraction. The results confirm the superiority of the proposed approach regarding classification accuracy by 2% to 6% improvement over other unsupervised feature learning methods. This paper introduced a representation learning method called Regularised Deep Convolution GAN (RDCGAN). Although Deep Convolution GAN (DCGAN)s achieve acceptable accuracy, they have still some forms of instability during training. RDCGAN sparked from Mode Regularized GANs gives evidence that using more feature maps not only can dramatically approve the model performance as a representation learning method for supervised tasks but also can stabilize mode in GANs. They also noticed that using reconstruction error can show a meaningful error curve corresponding to the generated images quality.

[4] Generative adversarial network (GAN) is a powerful generative model. However, it suffers from two key problems which are convergence instability and mode collapse. Recently, progressive growing of GANs for improving quality, stability and variation (PGGAN) is proposed to better solve these two problems. Although the performance of PGGAN is good on these two problems, it is still not satisfied on mode collapse problem. In this paper, they propose a new architecture based on PGGAN called D2PGGAN to better solve the mode collapse problem. The key idea consists of one generator and two different discriminators in PGGAN. With the fact that GAN is the analogy of a minimax game, the proposed architecture is as follows. The generator (G) aims to produce realistic looking samples to fool both of two discriminators. The first discriminator (D1) rewards high scores for samples from the data distribution, while the second one (D2) favour sample from the generator conversely. Specifically, a novel loss function is designed to optimize the proposed D2PGGAN. Extensive experiments on CIFAR-10 and CIFAR-100 datasets demonstrate that the proposed method is effective and obtains the highest inception scores compared with others state-of-the-art GANs. In this paper, they also project light towards the fact that some inception scores of the others GANs are different from the original papers in which they use chainer but not TensorFlow to obtain the inception scores. This is because more and more work of GAN using chainer, using chainer is convenient to evaluate the model. Because GAN will obtain different results in different times, they show that the average inception scores computed from 5 times run and the value is 8.59, this value is still higher than other GANs. To further showing that GAN is better, the images generated on CIFAR-10 which includes DCGAN, Minibatch discrimination (MD), D2GAN and Our GAN, it can be seen that the images generated by this GAN is better and more diversity. In this paper, they have chosen CIFAR – 100 Dataset for two reason. The first reason is that CIFAR-100 has the same inception model as CIFAR-10, which makes it convenient to evaluate the model with inception scores, while the other datasets cannot do that. The second reason is that CIFAR - 100 has more classes, which means that each class has fewer samples to use. Using fewer samples to generate images is also a meaningful work. In the quantitative results of this inception scores, it can be observed that their model yields the highest inception score compared with other state-of-the-art GANs. As the results of the CIFAR-10 dataset, all the others GANs are best-run results and obtained by themselves. The real data on the inception scores of CIFAR- 100 is 15.06 which is much higher than that of CIFAR-10 because CIFAR-100 has more classes than CIFAR-10. Because GAN will obtain different results in different times during training, we show the average inception scores computed from 5 times run and the value is 8.11, this value is higher than all of the other state-of-the-art GANs. In general, the inception scores of CIFAR-100 are lower than CIFAR- 10 because fewer samples of each classes is proposed. Their method still yields highest inception scores and this implies its superiority. Additionally, several samples generated by their proposed model and the other GANs on the CIFAR-100 dataset, the objects are becoming harder to recognize, but it can still be observed that their method generates better images with higher diversity.

[5] Image forensics aims to detect the manipulation of digital images. Currently, splicing detection, copy-move detection and image retouching detection are attracting significant attentions from researchers. However, image editing techniques develop over time. An emerging image editing technique is colorization, in which grayscale images are colorized with realistic colors. Unfortunately, this technique may also be intentionally applied to certain images to confound object recognition algorithms. To the best of their knowledge, no forensic technique has yet been invented to identify whether an image is colorized. They observed that, compared to natural images, colorized images, which are generated by three state-of-the-art methods, possess statistical differences for the hue and saturation channels. Besides, they also observe statistical inconsistencies in the dark and bright channels, because the colorization process will inevitably affect the dark and bright channel values. Based on their observations, i.e., potential traces in the hue, saturation, dark and bright channels, they propose two simple yet effective detection methods for fake colorized images: Histogram based Fake Colorized Image Detection (FCID-HIST) and Feature Encoding based Fake Colorized Image Detection (FCID-FE). Experimental results demonstrate that both proposed methods exhibit a decent performance against multiple state-of-the-art colorization approaches. In this paper, they aimed to address a new problem in the field of fake image detection: fake colorized image detection. They observed that fake colorized images and their corresponding natural images possess statistical differences in the hue, saturation, dark and bright channels. We proposed two simple yet effective schemes, FCID-HIST and FCID-FE, to resolve this detection problem. FCID-HIST exploits the most distinctive bins and total variations of the normalized histogram distributions and creates features for detection, while FCID-FE models the data samples with GMM and creates Fisher vectors for better utilizing the statistical differences.

[6] In this paper, they propose a reversible data hiding scheme that enables an adjustable amount of information to be embedded in JPEG images based on padding strategy. The proposed embedding algorithm only modifies, in a subtle manner, an adjustable number of zero-valued quantised DCT coefficients to embed the message. Hence, compared with a state-of-the-art based on histogram shifting, the proposed scheme has a relatively low distortion to the host images. In addition to this, they found that by representing the message in ternary instead of in binary, we can embed a greater amount of information while the level of distortion remains unchanged. Experimental results support that the proposed scheme can achieve better visual quality of the marked JPEG image than the histogram shifting based scheme. The

proposed scheme also outperforms this state-of-the-art in terms of the ease of implementation. In this paper, an easy-to-implement RDH scheme for JPEG images with an adjustable embedding capacity is proposed. Observed from the experimental results, the proposed method outperforms the state-of-the-art in terms of the visual quality since only a small number of zero valued coefficients are modified. However, the size of the compressed file increases to a slight extend since the modification on the zero-valued coefficients undermine the efficiency of RLE. In data hiding, it is essential to study the trade-offs between conflicting criteria and establish design principles that approach the theoretical limits. Furthermore, the research on RDH is expected to be developed for various host media and applications in a near future.

III.METHODOLOGY

In this paper we would like to write about a prototype of the model with respect to the survey done. Our project works on the principle of DeepFake, as we require a thousands of images irrespective of real or fake for the proper working as well as increase in the accuracy of our project. For this, we use the concept of DeepFake. After the collection of datasets then we pre-process the data with Error Level Analysis (ELA). Once we done that, we should split the data into train and test the training data will be feed through inside the Convolutional Neural Network (CNN) for train the model. Once it's getting trained, we use that trained model for evaluate the result for our test data. The illustration of the methodology is as shown in the below diagram.

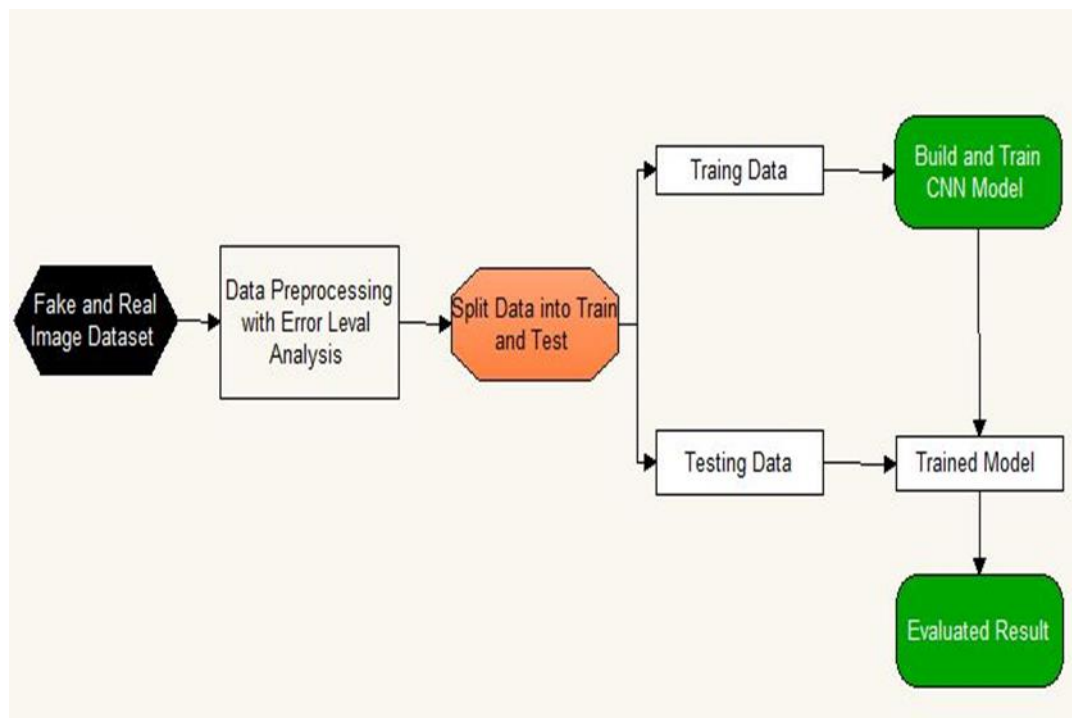


Fig. 1 Block diagram of the proposed system

IV.CONCLUSION

Based on the above study, the following conclusions were made:

- With the proposed study, It can be observed that fake images and their corresponding real images in the Error Level Analysis.
- The study aims to verify the authenticity of digital images without any prior knowledge of the original image.
- Mainly from this project is when we give new digital image to our system it's going to apply the ELA and CNN Algorithm's to that image then it will classify the given image is Fake or Real.

REFERENCES

- [1]. Yi-Lei Chen and Chiou-Ting Hsu, "Detecting Doubly Compressed Images Based on Quantization Noise Model and Image Restoration" IEEE International Workshop on Multimedia Signal Processing, Oct. 2009.
- [2]. Yuanman Li and Jiantao Zhou, "Fast and Effective Image Copy-Move Forgery Detection via Hierarchical Feature Point Matching" IEEE Transactions on Information Forensics and Security, vol.14, May 2019.
- [3]. Mehran Mehralian and Babak Karasfi, "RDCGAN: Unsupervised Representation Learning with Regularized Deep Convolutional Generative Adversarial Networks" 2018 9th Conference on Artificial Intelligence and Robotics and 2nd Asia-Pacific International Symposium, Dec 2018.
- [4]. Zhaoyu Zhang, Mengyan Li and Jun Yu, "D2PGGAN: Two Discriminators used in Progressive Growing of GANs" ICASSP 2019 - 2019 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP), May 2019.
- [5]. Yuanfang Guo, Xiaochun Cao, Wei Zhang and Rui Wang, "Fake Colorized Image Detection", IEEE Transactions on Information Forensics and Security, vol.13, Aug 2018.
- [6]. Ching-Chun Chang and Chang-Tsun Li, "Reversible Data Hiding in JPEG Images Based on Adjustable Padding", 2017 5th International Workshop on Biometrics and Forensics (IWBF), April 2017.
- [7]. H. Farid, "Image Forgery Detection," IEEE Signal Processing Magazine, 26(2), pp. 16-25, 2009.
- [8]. T. Bianchi and A. Piva, "Image Forgery Localization via Block-Grained Analysis of JPEG Artifacts," IEEE Trans. Inf. Forensics Security, vol. 7, no. 3, pp. 1003-1017, 2012.
- [9]. P. Ferrara, T. Bianchi, A. D. Rosa and A. Piva, "Image Forgery Localization via Fine-Grained Analysis of CFA Artifacts," IEEE Trans. Inf. Forensics Security, vol. 7, no. 5, pp. 1566-1577, 2012.
- [10]. H. Cao and A. C. Kot, "Accurate Detection of Demos using Regularity for Digital Image Forensics," IEEE Trans. Inf. Forensics Security, vol. 4, no. 4, pp. 899-910, 2009.
- [11]. Y. Li, J. Zhou, and A. Cheng, "SIFT key point removal via directed graph construction for color images," IEEE Trans. Inf. Forensics Security, vol. 12, no. 12, pp. 2971-2985, Dec. 2017
- [12]. D. Cozzolino, G. Poggi, and L. Verdoliva, "Efficient dense-field copy-move forgery detection," IEEE Trans. Inf. Forensics Security, vol. 10, no. 11, pp. 2284-2297, Nov. 2015.
- [13]. Z. Ni, Y. Q. Shi, N. Ansari, and S. Wei, "Reversible data hiding," IEEE Trans. Circuits Syst. Video Technol., vol. 16, no. 3, pp. 354-362, Mar. 2006.
- [14]. X. Li, B. Li, B. Yang, and T. Zeng, "General framework to histogram shifting-based reversible data hiding," IEEE Trans. Image Process., vol. 22, no. 6, pp. 2181-2191, Jun. 2013.
- [15]. T. Kalker and F. M. J. Willems, "Capacity bounds and constructions for reversible data-hiding," in Proc. Int. Conf. Digit. Signal Process., vol. 1, pp. 71-76, Jun. 2003.
- [16]. W. Zhang, B. Chen, and N. Yu, "Improving various reversible data hiding schemes via optimal codes for binary covers," IEEE Trans. Image Process., vol. 21, no. 6, pp. 2991-3003, Jun. 2012.
- [17]. X. Zhang, "Reversible data hiding with optimal value transfer," IEEE Trans. Multimedia, vol. 15, no. 2, pp. 316-325, Feb. 2013.
- [18]. W. Zhang, X. Hu, X. Li, and N. Yu, "Recursive histogram modification: Establishing equivalency between reversible data hiding and lossless data compression," IEEE Trans. Image Process., vol. 22, no. 7, pp. 2775-2785, Jul. 2013.
- [19]. J. Zheng, M. Liu, "A digital forgery image detection algorithm based on wavelet homomorphic filtering," Proc. of IWDW, LNCS vol. 5450, pp. 152-160, 2009.
- [20]. K. M. Carter, A. O. Hero, and R. Raich, "De-biasing for intrinsic dimension estimation," in Proc. Of the 14th Workshop on Statistical Signal Processing, 2007, pp. 601-605.
- [21]. Carter K. M., Raich R, and Hero A. O., "On local intrinsic dimension estimation and its applications," IEEE Trans. on Signal Processing, vol. 58, no. 2, pp. 650-663, 2010.
- [22]. A. Jessica Fridrich, B. David Soukal, and A. Jan Lukáš, "Detection of copy-move forgery in digital images," in Proceedings of Digital Forensic Research Workshop, 2003.
- [23]. AP Farid and AC Popescu, "Exposing digital forgeries by detecting duplicated image regions," Technical Report, TR2004-515, Department of Computer Science, Dartmouth College, Hanover, New Hampshire 2004.
- [24]. Weiqi Luo, Jiwei Huang, and Guoping Qiu, "Robust detection of region-duplication forgery in digital image," in Pattern Recognition, 2006. ICPR 2006. 18th International Conference on, 2006, pp. 746-749.
- [25]. Sevinc Bayram, Husrev Taha Sencar, and Nasir Memon, "An efficient and robust method for detecting copy-move forgery," in Acoustics, Speech and Signal Processing, 2009. ICASSP 2009. IEEE International Conference on, 2009, pp. 1053-1056.

BIOGRAPHY



Dr. Chanda.V. Reddy, completed PhD in wireless communication under VTU Belagavi, Karnataka, India. Is presently working as Professor and Head, Department of TCE, KSIT, Bangalore. Has 25 years of teaching experience with 11 papers published in reputed international journals. Presented 2 papers in national conferences one amongst it has won the award of best paper in national conference presented in RVCE Bangalore and presented and published 3 papers in international conferences.