

طراحی الگوریتم شناسایی تصاویر جعلی با استفاده از شبکه عصبی پیچشی

فاطمه زارع مهرجردی^{۱*}، علی محمد لطیف^۲، محسن سرداری زارچی^۳

۱- دکتری هوش مصنوعی و رباتیک، دانشگاه یزد - zarefateme41@yahoo.com

۲- دانشیار مهندسی کامپیوتر، دانشگاه یزد - alatif@yazd.ac.ir

۳- دانشیار مهندسی کامپیوتر، دانشگاه میبد - sardari@meybod.ac.ir

چکیده:

شناسایی تصاویر جعل یکی از موضوعهای چالش برانگیز در زمینه‌ی بینایی کامپیوتر است. تصاویر منبعی از اطلاعات هستند و در برخی از کاربردها به عنوان سند و شاهد مورد استفاده قرار می‌گیرند، از این رو اعتبار و صحت تصاویر مهم است. امروزه عمل جعل یا دست‌کاری تصویر با استفاده از ابزارهای ویرایش تصویر به راحتی انجام می‌شود. در دست‌کاری تصویر هدف، تغییر دادن مفهوم تصویر و در عین حال حفظ حداکثری یکپارچگی بافت و ساختار تصویر است. یکی از ساده‌ترین روش‌های دست‌کاری تصویر روش جعل کپی-انتقال است. در این روش بخشی از تصویر کپی و در مکان دیگری در همان تصویر منتقل می‌شود. این عمل با هدف پنهان‌سازی جزئیات خاصی از تصویر یا تکثیر جلوه‌های ویژه در تصویر صورت می‌گیرد. در این پژوهش، هدف طراحی شبکه یادگیری عمیق جهت شناسایی تصاویر جعلی از تصاویر سالم است و برای این منظور دو روش ارائه شده است. در روش اول یک شبکه با سه انشعاب مختلف متشکل از لایه‌های کانولوشن با سایز فیلترهای متفاوت طراحی شده است. این شبکه با تصاویر رنگی و تصاویر حاصل از تجزیه و تحلیل خطا، آموزش و مورد ارزیابی قرار گرفته است. در روش دوم یک شبکه با دو انشعاب متشکل از لایه‌های کانولوشن کاملاً مشابه اما با دو ورودی متفاوت، تصاویر رنگی و تصاویر حاصل از تجزیه و تحلیل خطا طراحی شده است. در هر دو روش انشعاب‌ها با هم ادغام شده و تصاویر به دو گروه جعل و سالم طبقه‌بندی می‌شوند. روش‌های پیشنهادی با استفاده از پایگاه داده MICC مورد ارزیابی قرار گرفته است و نتایج ارزیابی‌ها عملکرد رضایت‌بخش روش‌های پیشنهادی را نشان می‌دهد.

کلید واژه‌ها: یادگیری عمیق، جعل کپی-انتقال، تجزیه و تحلیل خطا، شبکه عصبی کانولوشنی.

Designing model to identify forgery images using convolutional neural network

Fatemeh Zare Mehrjardi¹, Ali Mohammad Latif², Mohsen Sardari Zarchi³

1- PhD in artificial intelligence and robotics, Yazd University - zarefateme41@yahoo.com

2- Associate Professor of Computer Engineering, Yazd University - alatif@yazd.ac.ir

3- Associate Professor of Computer Engineering, Meybod University - sardari@meybod.ac.ir

Abstract

Detecting forgery images is one of the most challenging topics in computer vision. Images are a source of information and are used as documents and evidence in various applications, so the validity and accuracy of images are crucial. Today, it is easy to manipulate an image using image editing tools. In image manipulation, the goal is to change the concept of the image while maintaining the maximum integrity of its texture and structure. Copy-move forgery is one of the simplest image manipulation methods. In this method, a part of the image is copied and pasted to another place in the same image with the purpose of hiding certain details of the image or adding special effects. The purpose of this research is to design a deep learning network that can identify forgery images from healthy ones. Two methods are presented for this purpose. In the first method, a network with three different branches consisting of convolution layers with different filter sizes is designed. This network has been trained and evaluated using color images and error level analysis images. In the second method, a network with two branches consisting of completely similar convolution layers but with two different inputs, color images and error level analysis images, is designed. In both methods, the branches are merged, and the images are classified into forgery and healthy groups. The proposed methods have been evaluated using

the MICC dataset, and the evaluation results show the satisfactory performance of the proposed methods.

Keywords: Deep Learning, Copy-move forgery, Error Analysis, Convolutional Neural Network.

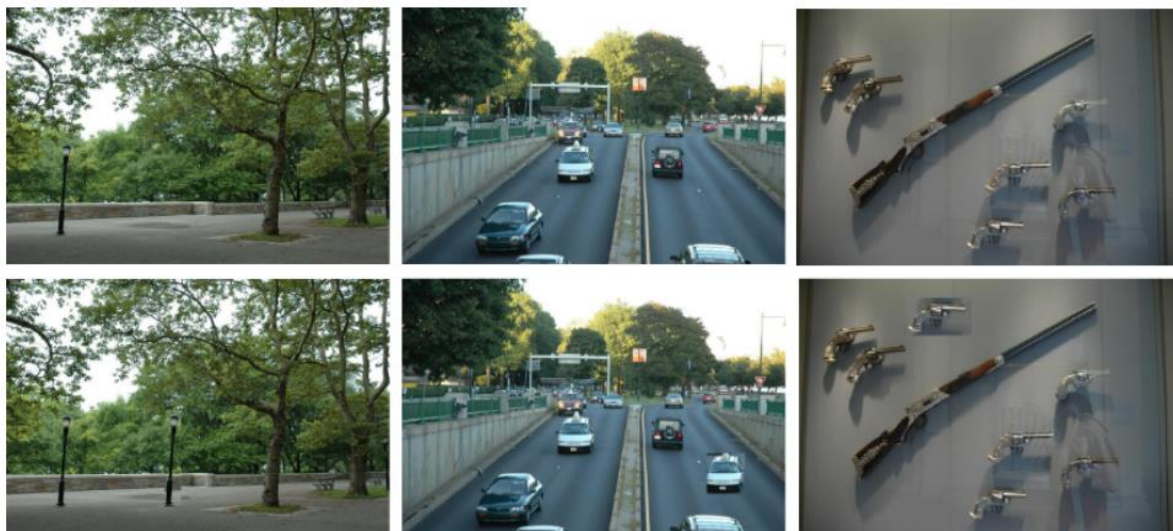
۱- مقدمه

تصویر به عنوان یکی از ابزارهای ارتباطی قوی و مهم بین انسان‌ها شناخته شده است. این یک حقیقت است که ارزش یک تصویر بیشتر از هزار کلمه است. با توسعه و گسترش دستگاه‌های دیجیتالی مانند دوربین و موبایل در هر لحظه از زمان و مکان تصویربرداری به راحتی امکان‌پذیر است. در بعضی از کاربردها تصاویر دیجیتال می‌توانند به عنوان سند مورد استفاده قرار گیرند. بنابراین، صحت و اعتبار تصویر اهمیت دارد. جعل تصویر به معنای دست‌کاری در تصویر دیجیتال است. امروزه با استفاده روزافزون از نرم افزارهای مدرن ویرایش تصویر مانند فتوشاپ دست‌کاری تصاویر دیجیتال به راحتی انجام می‌شود.

هدف دست‌کاری، پنهان کردن برخی از اطلاعات مفید یا اضافه کردن اطلاعات غلط یا اضافی به تصویر است. با جعل تصویر یکپارچگی و اصالت تصویر در ساختار و بافت تصویر دست‌کاری می‌شود. اگر تصویر مورد استفاده به عنوان سند دست‌کاری شود دیگر قابل اعتماد نخواهد بود. تصاویر دیجیتال به کمک روش‌های مختلفی جعل می‌شوند. رویکردهای موجود برای تشخیص جعل تصاویر دیجیتال را می‌توان به طور کلی به دو دسته‌ی عمده رویکرد فعال و رویکرد غیرفعال طبقه‌بندی کرد [۱].

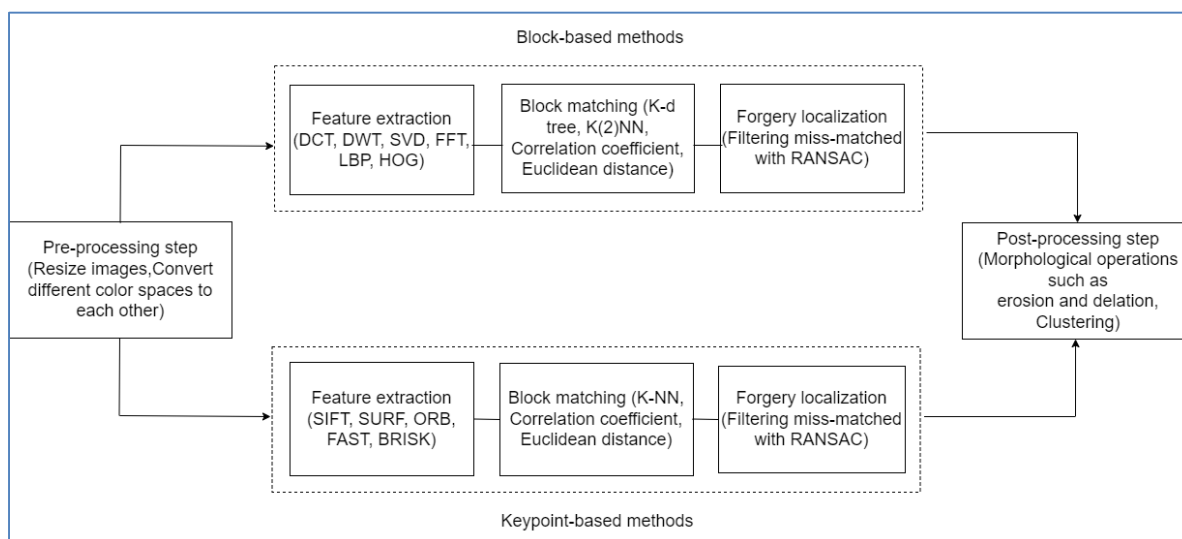
در رویکرد فعال، ابتدا بر روی تصاویر پیش پردازش اعمال شده، سپس، با تعبیه کردن اطلاعات درون تصاویر اصلی عمل دست‌کاری انجام می‌شود. امضای دیجیتال و واترمارکینگ را می‌توان به عنوان مثال‌هایی از رویکرد فعال نام برد. در رویکرد غیرفعال که به عنوان رویکرد کور شناخته می‌شود، نیازی به پیش پردازش تصویر نیست. رویکرد غیرفعال با تحلیل محتوا و ساختار تصویر، جعلی بودن یا نبودن تصویر را تشخیص می‌دهد و نا سازگاری‌ها را کشف می‌کند. جعل کپی - انتقال^[۲] اتصال تصویر^[۳]، روتوش تصویر^[۴] حذف شی^[۵]، جمله مثال‌های رویکرد غیرفعال است [۲].

روش کپی - انتقال یکی از ساده‌ترین و رایج‌ترین روش‌های دست‌کاری تصاویر است. این روش بخشی از تصویر را کپی و در محل دیگر از همان تصویر انتقال می‌دهد. انگیزه‌ی چنین جعلی، پنهان کردن بخشی از تصویر، تأکید بر بخشی از تصویر و درج اطلاعات غلط در تصویر است. از آن جایی که بخش کپی شده متعلق به همان تصویر است، از لحاظ پویایی و رنگ با بقیه تصویر سازگاری دارد. همراه با عمل کپی و انتقال تبدیل‌های هندسی مانند چرخش^[۶] و مقیاس‌بندی^[۷] عملیات پس پردازش مانند محو کردن، تغییر روشنایی، فشردگی و افزودن نویز به تصویر نیز انجام می‌شود تا ناحیه دست‌کاری شده برای چشم انسان به راحتی قابل تشخیص نباشد [۱ و ۲]. شکل ۱ مثالی از جعل کپی - انتقال را نشان می‌دهد.



شکل (۱): تعدادی از تصاویر موجود در پایگاه داده MICC (سطر اول تصاویر سالم و سطر دوم تصاویر جعلی با روش کپی - انتقال) [۳].

با توجه به مطالعه‌های انجام شده، الگوریتم‌های تشخیص جعل کپی - انتقال را می‌توان به دو گروه اصل‌سنجی و مُبتنی بر یادگیری عمیق تقسیم‌بندی کرد. الگوریتم‌های تشخیص سنتی خود به دو روش بر پایه‌ی بلوک‌بندی و یافتن نقاط کلیدی^۱ تصویر تقسیم می‌شوند. شکل ۲ مراحل الگوریتم‌های تشخیص سنتی را نشان می‌دهد.



شکل (۲): مراحل تشخیص جعل به روش سنتی [۲].

در روش بر پایه‌ی بلوک‌بندی ابتدا تصویر به نواحی مستطیل یا دایره‌ای شکل هم‌پوشان و غیرهم‌پوشان تقسیم شده و برای هر کدام از این نواحی یک بردار ویژگی محاسبه می‌شود. در مرحله بعد بردارهای ویژگی مشابه با هم انطباق می‌یابند و حمله کپی - انتقال تشخیص داده می‌شود. عیب این روش پیچیدگی محاسباتی بالا و عملکرد ضعیف در برابر تغییرات هندسی و عملیات پس‌پردازش است [۴].

در روش‌های مُبتنی بر نقاط کلیدی ابتدا ویژگی‌های نواحی با آنتروپی بالا به‌دست می‌آیند و سپس، ویژگی‌های مشابه تطبیق می‌یابند. در روش مُبتنی بر نقاط کلیدی، ویژگی با استفاده از روش‌های مختلف مانند SIFT، SURF بدون تقسیم‌بندی تصویر استخراج می‌شود. نقاط ویژگی با استفاده از رویکردهای مختلف مانند خوشه‌بندی، فاصله اقلیدسی با یکدیگر

تطابق داده شده و جعل مشخص می‌شود. این روش دارای محاسبات کم و مقاومت مناسب در برابر تبدیل‌های هندسی است [۵].

روش یادگیری عمیق امروزه در بیشتر مسائل وارد شده است و نتایج خوبی را به دست آورده است. الگوریتم‌های یادگیری عمیق زیرمجموعه‌ای از الگوریتم‌های یادگیری ماشین هستند که هدف آنها کشف چندین سطح از بازنمودهای (نمایش‌های) توزیع شده از داده ورودی است. برخلاف دو روش بر پایه بلوک‌بندی و نقاط کلیدی، روش یادگیری عمیق تمام مراحل را به صورت خودکار انجام می‌دهد و ویژگی‌های سطح بالا را از روی داده‌ها استخراج می‌کند و از استخراج ویژگی به صورت دستی جلوگیری می‌کند. یکی از چالش‌های روش یادگیری عمیق نیاز به داده زیاد برای آموزش است. استفاده از شبکه‌های از پیش آموزش داده‌شده با پایگاه داده‌های بزرگ از جمله ImageNet، روش انتقال دانش و روش‌های تقویت داده برای افزایش تعداد داده‌ها برای حل مشکل کمبود داده ارائه شده‌اند [۱ و ۲].

۲- مروری بر کارهای گذشته

در سال‌های اخیر مطالعه‌های زیادی برای تشخیص جعل کپی - انتقال با استفاده از سه رویکرد براساس بلوک‌بندی، مبتنی بر نقاط کلیدی و یادگیری عمیق انجام شده است. خلاصه‌ای از برخی از این مطالعه‌ها در این بخش آورده شده است. برای مثال، یک روش مبتنی بر بلوک‌بندی با استفاده از ترکیب دو الگوریتم DCT و PCA در سال ۲۰۱۶م ارائه شده است. برای این منظور ابتدا در مرحله پیش‌پردازش، تصاویر رنگی به تصاویر سطح خاکستری تبدیل شده است. سپس، تصاویر به بلوک‌های مربعی شکل و به صورت هم‌پوشان تقسیم‌بندی شده‌اند و از هر بلوک با استفاده از الگوریتم DCT یک بردار ویژگی استخراج شده است. بردارهای ویژگی استخراج شده دارای ابعاد فضای ویژگی بالایی هستند، جهت کاهش ابعاد فضای ویژگی از الگوریتم PCA استفاده شده است. در نهایت برای تشخیص جفت بلوک مشابه و تشخیص ناحیه جعل، از مقایسه فاصله اقلیدسی بین تمام جفت بلوک‌ها استفاده شده است [۶].

در مطالعه‌های دیگر جهت تشخیص جعل کپی - انتقال از روش مبتنی بر بلوک‌بندی و ترکیب دو الگوریتم DCT و SVD استفاده شده است. ابتدا در مرحله پیش‌پردازش تصاویر رنگی به تصاویر خاکستری تبدیل شده و تصاویر به بلوک‌های هم‌پوشان تقسیم‌بندی شده‌اند. در مرحله استخراج ویژگی از ترکیب دو الگوریتم DCT و SVD جهت به دست آوردن بردارهای ویژگی کوچک و منحصر به فرد استفاده شده است. سپس، در مرحله تطابق ویژگی و یافتن جفت بلوک مشابه از مقایسه جفت بلوک‌ها با استفاده از روش فاصله اقلیدسی استفاده شده است [۷].

در روش مبتنی بر نقاط کلیدی، برخلاف روش مبتنی بر بلوک‌بندی نیازی به تقسیم‌بندی تصاویر نیست و با استخراج نقاط کلیدی از نواحی با آنتروپی بالا، عمل تشخیص جعل انجام می‌شود، از این رو نسبت به روش بلوک‌بندی سریع‌تر است. برای مثال یک روش تشخیص جعل با نقاط کلیدی و با استفاده از الگوریتم استخراج‌کننده نقاط کلیدی به نام SIFT توسط Alberry و همکاران ارائه شده است. از این‌رو در مرحله استخراج ویژگی، بدون تقسیم‌بندی تصویر و با استفاده از الگوریتم SIFT نقاط کلیدی تصاویر کشف شده است. سپس، در مرحله تطابق ویژگی از الگوریتم خوشه‌بندی C-means جهت شناسایی نقاط کلیدی مشابه استفاده شده است [۸].

یکی از مشکلات روش نقاط کلیدی عملکرد ضعیف این روش در تشخیص جعل نواحی هموار و کوچک به‌خاطر کمبود نقاط کلیدی کافی است. برای حل این مشکل محققان معمولاً از ترکیب تعدادی الگوریتم استخراج‌کننده نقاط کلیدی استفاده می‌کنند. برای مثال یک روش تشخیص جعل مبتنی بر نقاط کلیدی با استفاده از ترکیب دو الگوریتم استخراج‌کننده نقاط کلیدی SURF و A-KAZE ارائه شده است. بعد از استخراج نقاط کلیدی با استفاده از الگوریتم فاصله اقلیدسی، شباهت جفت نقاط کلیدی محاسبه و مرتب شده است. سپس، با استفاده از الگوریتم دو نزدیک‌ترین همسایگی نقاط کلیدی مشابه کشف شده است [۹].

برخی محققان از ترکیب دو روش نقاط کلیدی و روش بلوک‌بندی جهت تشخیص جعل استفاده می‌کنند. یک نمونه از روش ترکیبی توسط Narayanan و همکاران ارائه شده است. بدین ترتیب ابتدا تصاویر به بلوک‌های مستطیل شکل و هم‌پوشان

تقسیم‌بندی شده و نقاط کلیدی هر بلوک با استفاده از الگوریتم استخراج‌کننده نقاط کلیدی SIFT کشف شده است. در مرحله تطابق ویژگی از الگوریتم فاصله اقلیدسی جهت یافتن نقاط کلیدی مشابه هر بلوک استفاده شده است. در پایان اگر تعداد نقاط کلیدی تطابق‌یافته از سطح آستانه بیشتر باشد، جفت بلوک مشابه شناسایی شده است [۱۰].

بر خلاف دو روش نقاط کلیدی و مبتنی بر بلوک‌بندی که شامل مراحل برای استخراج ویژگی و تطابق ویژگی جهت تشخیص جعل هستند، روش یادگیری عمیق تمام مراحل این فرآیند را به صورت یک‌جا و به صورت خودکار انجام می‌دهد، اما جهت آموزش شبکه‌های یادگیری عمیق، داده زیادی مورد نیاز است. برای مثال یک روش مبتنی بر یادگیری عمیق جهت شناسایی تصاویر جعلی توسط Elaskily و همکاران ارائه شده است. در این روش به دلیل کمبود داده در حوزه جعل از شبکه‌های از پیش آموزش دیده شده با پایگاه داده بزرگ imagenet مفهوم یادگیری انتقالی استفاده شده و پارامترهای شبکه با داده‌های کم موجود تنظیم شده است. در نهایت تصاویر در دو دسته تصاویر اصلی و تصاویر جعل طبقه‌بندی شده‌اند [۱۱].

در مطالعه‌ای دیگر از شبکه یادگیری عمیق از پیش‌آموزش داده شده AlexNet جهت تشخیص جعل استفاده شده است. در این روش از شبکه از پیش‌آموزش داده شده AlexNet به عنوان یک واحد استخراج‌کننده ویژگی^[۱۲] استفاده شده است و از تمام تصاویر سالم و جعل یک بردار ویژگی استخراج شده است. در نهایت این بردارهای ویژگی به طبقه‌بند ماشین بردار پشتیبان^[۱۳] داده شده و تصاویر در دو دسته سالم و جعل دسته‌بندی شده‌اند [۱۲].

در این پژوهش هدف تشخیص تصاویر جعل با استفاده از روش یادگیری عمیق است. برای این منظور شبکه‌ی عصبی کانولوشن با بیش از یک انشعاب، با لایه‌ها و ورودی‌های مختلف و با استفاده از مفهوم API عملکردی طراحی شده است که در ادامه آورده شده است.

۳- مفاهیم پایه

۳-۱- شبکه عصبی کانولوشن

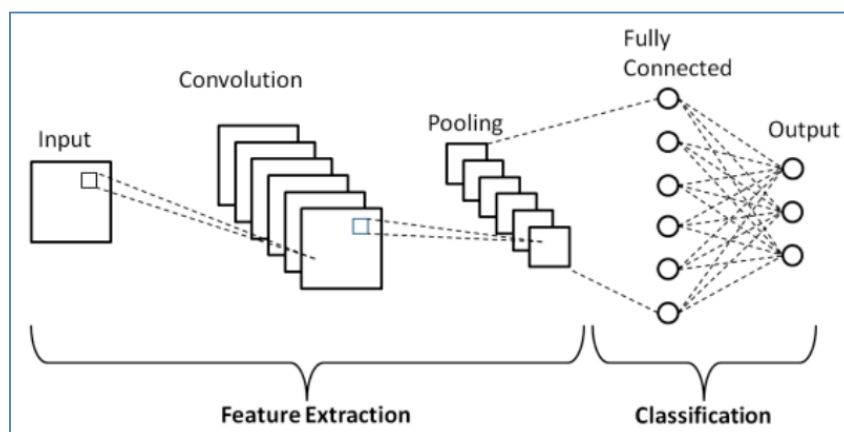
شبکه‌های عصبی کانولوشن یکی از مهم‌ترین روش‌های یادگیری عمیق در کاربردهای مختلف بینایی کامپیوتر هستند. شبکه عصبی کانولوشن نوعی شبکه عصبی مصنوعی است که از پرسپترون چندگانه استفاده می‌کند. یکی از مزایای استفاده از شبکه عصبی کانولوشن، استفاده از انسجام مکانی محلی در تصاویر ورودی است که به آنها اجازه می‌دهد تا با به اشتراک گذاشتن برخی پارامترها، وزن کمتری داشته باشند. این فرایند از نظر حافظه و پیچیدگی کارآمد است. به طور کلی، یک شبکه عصبی کانولوشن از سه لایه اصلی لایه کانولوشن^[۱۴]، لایه پولینگ^[۱۵] و لایه کاملاً متصل^[۱۶] تشکیل می‌شود و هر لایه وظیفه خاصی را بر عهده دارند [۱۳].

لایه کانولوشن به عنوان اصلی‌ترین لایه در این دسته از شبکه‌های عصبی محسوب می‌شود به همین دلیل، آنها را به عنوان شبکه‌های عصبی کانولوشن می‌شناسند. وظیفه‌ی این لایه استخراج ویژگی‌ها است. لایه‌های اول معمولاً ویژگی‌های سطح پایین (مانند لبه‌ها، خط‌ها و گوشه‌ها) و لایه‌های بالاتر ویژگی‌های سطح بالا (مانند قواعد و شکل‌ها) را استخراج می‌کنند. این لایه عملیات کانولوشن را بر روی داده ورودی اعمال می‌کند و خروجی‌هایی به نام نقشه ویژگی^[۱۷] از این لایه به دست می‌آید. در نتیجه، تمامی نرون‌ها در یک نقشه ویژگی، مجموعه‌ای از وزن‌ها و بایاس‌های مشابه و مشترکی دارند. این اشتراک وزن‌ها باعث کاهش تعداد پارامترهای مورد نیاز برای آموزش می‌گردد. در شبکه‌های کانولوشن با استفاده از فیلتر، اتصالات به صورت نواحی کوچک و محلی صورت می‌گیرد. به بیان دیگر هر نرون در نخستین لایه مخفی به ناحیه کوچکی از نرون‌های ورودی متصل است. بر روی هر تصویر می‌توان چندین نوع فیلتر اعمال کرد که در نتیجه به تعداد فیلترها، همان تعداد عمق ایجاد می‌شود. رایج‌ترین شکل استفاده از این لایه به صورت لایه‌ای با فیلترهایی به اندازه سائز 3×3 و اندازه گام^[۱۸] مقدار ۱ است.

در یک شبکه عصبی کانولوشن به طور معمول پس از هر لایه کانولوشن یک لایه پولینگ قرار می‌گیرد. این لایه از آن جهت اهمیت دارد که باعث کاهش تعداد پارامترها می‌شود. این لایه بر روی هر برش از لایه‌های عمق ورودی اعمال می‌شود و اندازه‌ی آن را تغییر می‌دهد. دو تابع عملکردی معروف این لایه، پولینگ بیشینه و میانگین‌نامدارند. طریقه عمل پولینگ بیشینه بدین صورت است که در هر پنجره، بزرگ‌ترین پیکسل را به خروجی می‌فرستد. به دلیل اعمال عملیات بر روی تمامی برش‌ها،

عمق خروجی همان عمق ورودی به لایه پولینگ است و تغییری نمی‌کند. به علت برش‌ها، عمق خروجی همان عمق ورودی به لایه پولینگ است و تغییری نمی‌کند. رایج‌ترین شکل استفاده از این لایه به صورت لایه‌ای با فیلترهایی به اندازه 2×2 و اندازه گام با مقدار ۲ است.

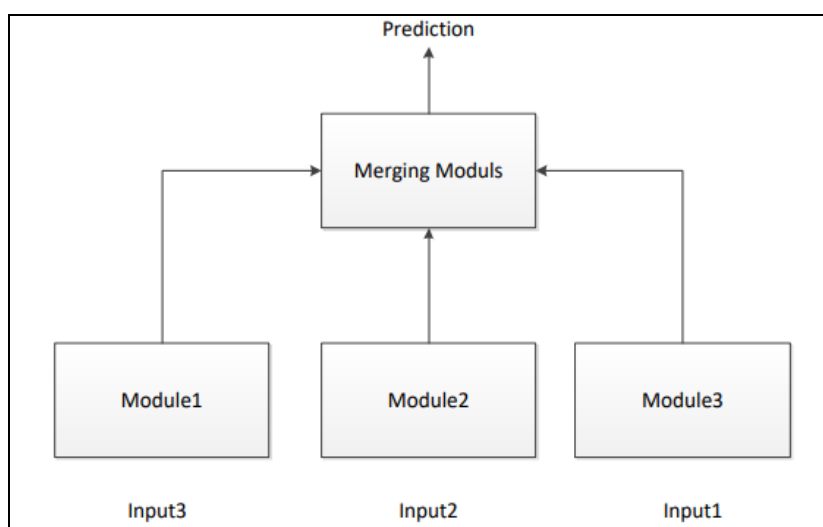
در انتهای یک شبکه عصبی کانولوشن، خروجی آخرین لایه پولینگ به عنوان ورودی به لایه کاملاً متصل عمل می‌کند. یک یا چند مورد از این لایه‌ها وجود دارد. لایه کاملاً متصل به این معنی است که هر گره در لایه اول به هر گره در لایه دوم وصل می‌شود. در واقع این لایه به عنوان لایه پایانی و لایه امتیازدهی برای خروجی سیستم در نظر گرفته می‌شود [۱۳]. شکل ۳ ساختار شبکه عصبی کانولوشن را نشان می‌دهد.



شکل (۳): ساختار شبکه عصبی کانولوشن [۱۳].

۳-۲- مفهوم API عملکردی

معمولاً معماری‌های یادگیری عمیق با استفاده از روش مدل ترتیبی ساخته می‌شوند. به این معنی که لایه‌ها همانند یک پشته خطی در کنار هم قرار دارند و هر لایه یک ورودی و یک خروجی دارد. گاهی در مسائل مختلف شبکه، ممکن است به تعدادی ورودی مستقل نیاز داشته باشد، برخی از شبکه‌ها باید بیش از یک خروجی داشته باشند و یا دارای انشعابات داخلی در بین لایه‌ها باشند، این موارد با مفهوم API شناخته می‌شوند [۱۴]. شکل ۴ طرز کار مفهوم API عملکردی را نمایش می‌دهد.



شکل (۴): طرز کار مفهوم API عملکردی.

۳-۳- تجزیه و تحلیل سطح خطا

تجزیه و تحلیل سطح خطا^[۳] مکان شناسایی مناطقی در یک تصویر که در سطوح مختلف فشرده سازی هستند، را نشان می‌دهد. در یک تصویر JPEG، کل تصویر باید در یک سطح از فشرده سازی باشد. اگر بخشی از تصویر دارای تفاوت قابل توجهی باشد، نشان‌دهنده یک تغییر دیجیتال است [۱۵]. ELA تفاوت در نرخ فشرده سازی JPEG را برجسته می‌کند. مناطق با رنگ‌بندی یکنواخت، مانند آسمان آبی یا دیوار سفید نسبت به لبه‌های با کنتراست بالا نتیجه ELA کمی خواهند داشت. بنابراین، اگر در یک تصویر جعل رخ داده باشد، نرخ فشرده سازی ناحیه جعل متفاوت خواهد بود، بنابراین، از ELA برای تشخیص جعل می‌توان استفاده کرد.

۳-۴- پایگاه داده MICC

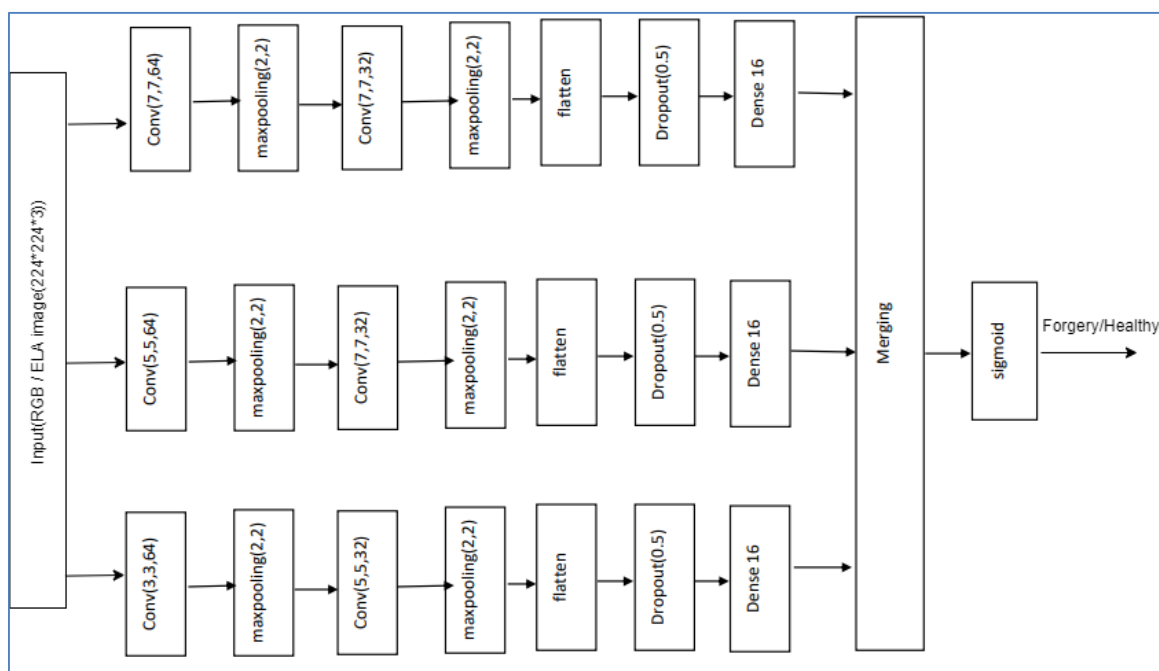
پایگاه داده MICC یکی از معروف‌ترین و قدیمی‌ترین پایگاه داده‌های موجود در زمینه تشخیص جعل است. این پایگاه داده دارای چهار زیرمجموعه با نام‌های MICC-F220، MICC-F600، MICC-F2000 و MICC-F8multi است. تصاویر موجود در این پایگاه داده با روش جعل کپی-انتقال به وجود آمده است. همراه با عمل جعل برخی از تبدیل‌های هندسی از جمله چرخش و مقیاس‌بندی بر روی ناحیه جعل اعمال شده است. پایگاه داده MICC دارای معایبی نیز هست که در ادامه آورده شده است [۳]:

- ۱- در این پایگاه داده تصاویر جعلی هیچ عملیات پس‌پردازشی ندارند.
- ۲- تصاویر جعلی موجود در سه زیرمجموعه MICC-F220، MICC-F2000 و MICC-F8multi بدون ماسک حقیقی است که نشان‌دهنده ناحیه جعل است.
- ۳- اندازه‌ی تصاویر در این پایگاه داده بزرگ است و برای الگوریتم‌های مبتنی بر بلوک‌بندی زمان‌بر است، این پایگاه داده برای الگوریتم‌های نقاز کلیدی مناسب است
- ۴- تصاویر سالم و جعل در این پایگاه داده محدود و نامتعادل است و تعداد تصاویر جعل کمتر است [۳].

۴- روش پیشنهادی

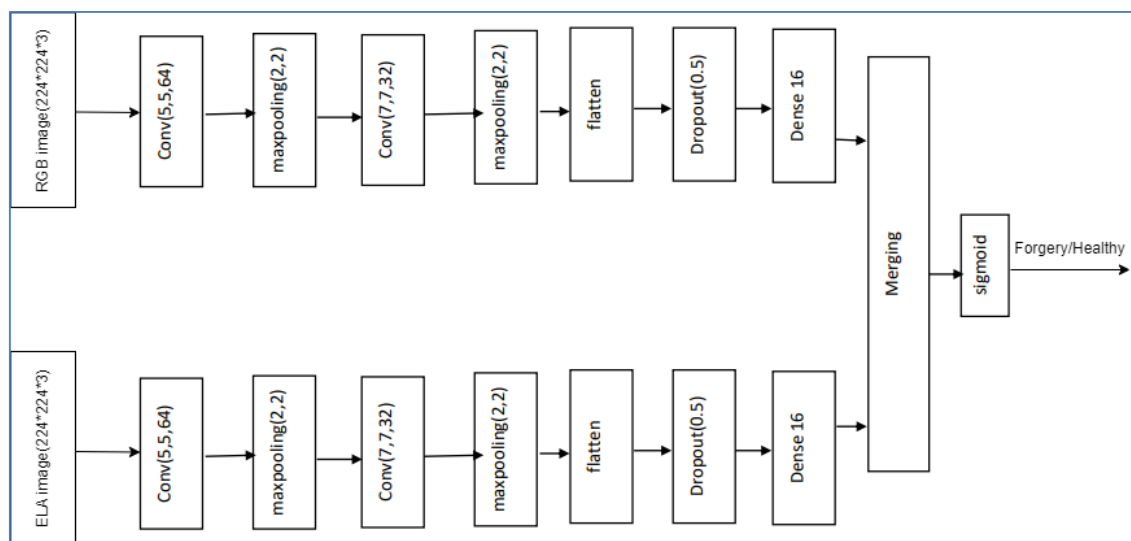
در این پژوهش با استفاده از مطالب مطرح شده در قسمت مفاهیم پایه دو روش متفاوت بررسی شده است.

روش اول: در این حالت از مفهوم API عملکردی استفاده شده و یک معماری کم عمق طراحی شده است که دارای سه انشعاب با تعداد لایه‌های یکسان، دو لایه کانولوشن، دو لایه پولینگ بیشینه و دو لایه کاملاً متصل است. از آنجا که اطلاعات مفید برای تشخیص جعل در لایه‌های ابتدایی شبکه موجود است از طراحی شبکه با تعداد لایه‌های بیشتر خودداری شده است. هر سه انشعاب از لحاظ تعداد لایه‌ها و تعداد فیلترها برابر هستند، ولی اندازه فیلترها متفاوت انتخاب شده است. در برخی از تصاویر، عمل جعل همراه با عملیات مقیاس انجام می‌شود انتخاب اندازه فیلترها با مقادیر ۳، ۵ و ۷ در سه انشعاب می‌تواند در تشخیص جعل همراه با عملیات مقیاس نقش مهمی را داشته باشد. در ادامه خروجی سه انشعاب با هم ادغام شده و پیش‌بینی نهایی با لایه امتیازدهی با یک نرون و تابع سیگموید انجام می‌شود. شبکه طراحی شده طی ۲۵ تکرار هم با تصاویر رنگی و هم با تصاویر حاصل از ELA آموزش داده می‌شود. شکل ۵ روندنمای رویکرد اول را نشان می‌دهد.



شکل (۵): روندنمای روش پیشنهادی اول.

روش دوم: در این حالت یک معماری دارای دو انشعاب با تعداد لایه‌ها، تعداد فیلترها و اندازه فیلتر یکسان، ولی با ورودی‌های متفاوت طراحی شده است. مانند رویکرد اول نقشه‌های ویژگی دو انشعاب با هم ادغام شده و با تابع سیگموید پیش‌بینی نهایی انجام می‌شود. شبکه طراحی شده طی ۲۵ بار تکرار با ورودی‌های متفاوت تصاویر رنگی و تصاویر حاصل از ELA آموزش داده می‌شود. شکل ۶ روند رویکرد دوم را نشان می‌دهد.



شکل (۶): روندنمای روش پیشنهادی دوم.

۵- ارزیابی نتایج

برای ارزیابی روش پیشنهادی از تصاویر سالم و جعل سه زیر مجموعه MICC-F2000، MICC-F600 و MICC-F220 استفاده شده است. از آن جا که مجموع تصاویر سالم با مجموع تصاویر جعل برابر نیست و در فرآیند آموزش احتمال سوگیری وجود دارد، به تعداد مساوی از هر کدام از مجموعه تصاویر جعل و سالم تصویر انتخاب شده است. برای افزایش تعداد تصاویر برای آموزش شبکه از روش تقویت داده استفاده شده است. در این روش با انجام عملیات انتقالی متفاوت مانند چرخش، برش دادن و بزرگ‌نمایی بر روی تصاویر موجود در پایگاه داده، تعداد تصاویر را افزایش می‌دهند.

برای ارزیابی روش اول یکبار تصاویر رنگی و بار دیگر تصاویر حاصل از ELA به شبکه که دارای سه انشعاب متفاوت است داده می شود. تعداد لایه ها و فیلترهای هر انشعاب برابر، ولی در اندازه فیلترها متفاوت هستند. در ادامه خروجی سه انشعاب با هم ادغام شده و در نهایت تصویر در یکی از دو دسته جعل و سالم قرار داده می شود. فرآیند آموزش شبکه پیشنهادی ۲۵ تکرار می شود. جدول ۱ میزان دقت و تابع زیان نتایج روش اول را نشان می دهد.

جدول (۱): نتایج ارزیابی روش پیشنهادی اول.

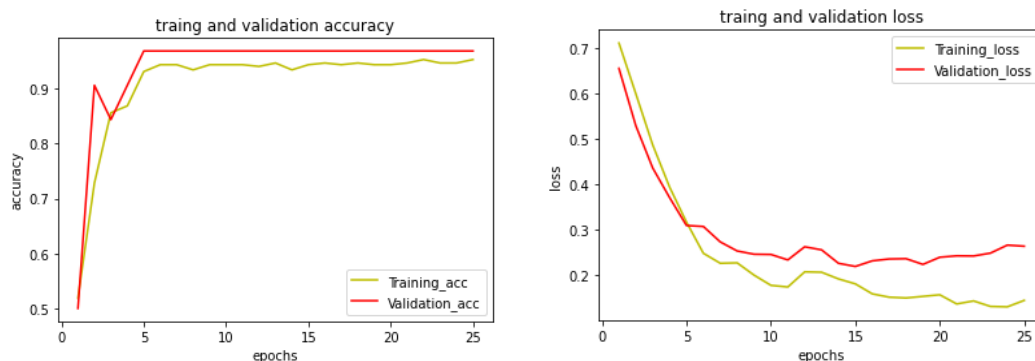
نوع تصاویر	معیار	Acc	Loss
رنگی	Train	۹۵/۳۱	۰/۱۴۵۲
	Test	۹۶/۸۸	۰/۲۱۹۷
ELA	Train	۹۵/۸۷	۰/۱۸۵۳
	Test	۹۶/۳۹	۰/۲۱۹۲

در روش دوم شبکه طراحی شده دارای دو انشعاب مشابه است و تعداد لایه ها، تعداد فیلترها و اندازه ی آنها مشابه در نظر گرفته شده است. برای ارزیابی این رویکرد دو ورودی متفاوت تصاویر رنگی و تصاویر حاصل از ELA به شبکه داده می شود. نقشه های ویژگی دو انشعاب با هم ادغام شده و با تابع سیگموئید تصاویر طبقه بندی می شوند. شبکه طی ۲۵ تکرار آموزش داده می شود. جدول ۲ میزان دقت و تابع زیان این رویکرد را نشان می دهد.

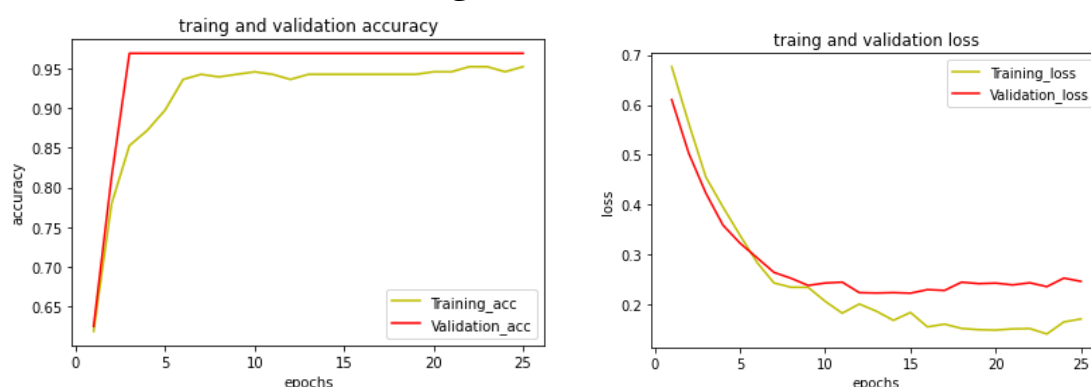
جدول (۲): نتایج ارزیابی روش پیشنهادی دوم.

	Acc	Loss
Train	۹۵/۳۰	۰/۱۴۲۲
Test	۹۶/۹۱	۰/۲۱۴۵

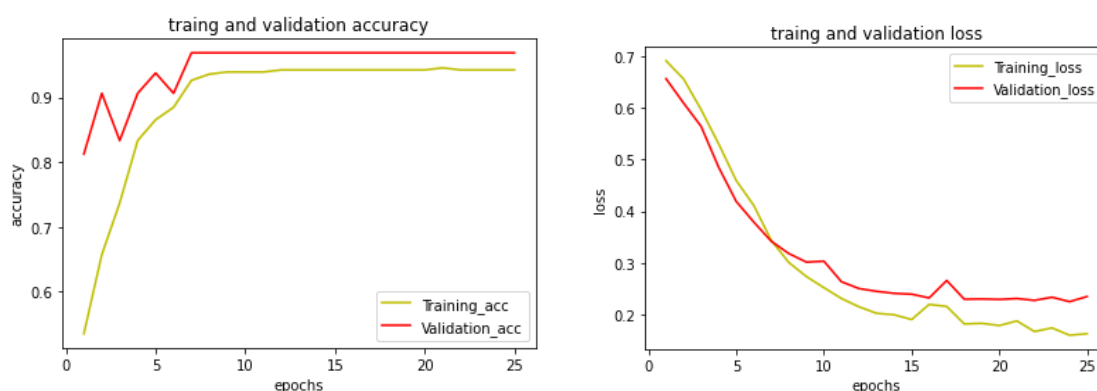
لازم به ذکر است برای هر دو روش پیشنهادی از بهینه ساز Adam با نرخ یادگیری ۰,۰۰۰۱ استفاده شده است. نمودار دقت و تابع زیان این دو روش بر روی تصاویر رنگی و تصاویر ELA در شکل ۷ آورده شده است.



روش اول (تصاویر رنگی)



روش اول (تصاویر ELA)



روش دوم

شکل (۷): نمودار دقت و تابع زیان دو روش پیشنهادی بر روی تصاویر رنگی و ELA.

برای مقایسه‌ی روش‌های پیشنهادی با روش‌های دیگران دو روش دیگر بررسی شده است. در مقاله [۱۲] از شبکه از پیش‌آموزش دیده‌شده AlexNet به عنوان استخراج‌کننده‌ی ویژگی استفاده شده است. از هر تصویر یک بردار ویژگی استخراج شده و سپس، با استفاده از الگوریتم ماشین بردار پشتیبان تصاویر در دو دسته سالم و جعل طبقه‌بندی شدند. در مقاله دیگر [۱۶] از شبکه از پیش‌آموزش داده‌شده MobileNetV2 و مفهوم انتقال دانش برای این منظور استفاده شده است. مقایسه در جدول ۳ آورده شده است.

جدول (۳): مقایسه روش پیشنهادی با کارهای دیگران.

روش	Accuracy
مقاله [۱۲]	۹۴/۴۶

مقاله [۱۶]	۹۳/۳۸
روش اول (تصاویر ELA)	۹۶/۳۹
روش اول (تصاویر رنگی)	۹۶/۸۸
روش دوم	۹۶/۹۱

۶- نتیجه‌گیری

تصاویر در زندگی امروزی انسان‌ها نقش مهمی دارد و از آن در برخی کاربردها به عنوان سند استفاده می‌شود. تصاویر با روش‌ها و ابزارهای مختلف و به سادگی می‌توانند دست‌کاری شوند، و اعتبار خود را از دست دهند. بنابراین، تشخیص تصویر جعل یک موضوع مهم در دنیای امروز به حساب می‌آید تا بتوان از حوادث ناگوار بعدی جلوگیری کرد. روش‌های سنتی موجود از قبیل روش بلوک‌بندی و نقاط کلیدی برای تشخیص جعل، دارای مشکلاتی است. هر دو روش تنها قادر به تشخیص تصاویر جعل از نوع کپی-انتقال هستند. روش بلوک‌بندی دارای محاسبات زیاد و ناکارآمدی در برابر تبدیلات هندسی است. روش نقاط کلیدی اگر چه در برابر تبدیلات هندسی مقاوم است، ولی وابستگی به نقاط کلیدی یکی از ضعف‌های این روش است. در سال‌های اخیر از روش‌های یادگیری عمیق در همه‌ی زمینه‌ها از جمله تشخیص جعل استفاده می‌شود. در این پژوهش از شبکه عصبی کانولوشن و مفهوم API عملکردی و تصاویر ELA استفاده شده و شبکه‌هایی با دو و سه انشعاب با ورودی‌های متفاوت و ساینز فیلترهای متفاوت برای غلبه بر تبدیل‌های هندسی و عملیات پس‌پردازش ارائه شده است. نتیجه‌ی ارزیابی‌ها و مقایسه با کارهای دیگران از عملکرد رضایت‌بخش، روش پیشنهادی است.

منابع:

- [1] Mehrjardi, F. Z., Latif, A. M., and Zarchi, M. S., "Copy-move forgery detection and localization using deep-learning," *International Journal of Pattern Recognition and Artificial Intelligence*, Vol. 37, No. 9, pp. 1-21, 2023.
- [2] Mehrjardi, F. Z., Latif, A. M., Zarchi, M. S., and Sheikhpour, R., "A survey on deep learning-based image forgery detection," *Pattern Recognition*, Vol. 144, pp. 1-31, 2023.
- [3] Amerini, I., Ballan, L., Caldelli, R., Del Bimbo, A., Del Tongo, L., and Serra, G., "Copy-move forgery detection and localization by means of robust clustering with J-Linkage," *Signal Processing: Image Communication*, Vol. 28, No. 6, pp. 659-669, 2013.
- [4] Mohamadian, Z., and Pouyan, A. A., "Detection of duplication forgery in digital images in uniform and non-uniform regions," In *15th International Conference on Computer Modelling and Simulation*, IEEE pp. 455-460, 2013.
- [5] Mahmood, T., Nawaz, T., Ashraf, R., Shah, M., Khan, Z., Irtaza, A., and Mehmood, Z., "A survey on block-based copy move image forgery detection techniques," In *International Conference on Emerging Technologies (ICET)*, IEEE, pp. 1-6, 2015.
- [6] Mahmood, T., Nawaz, T., Irtaza, A., Ashraf, R., Shah, M., and Mahmood, M. T., "Copy-move forgery detection technique for forensic analysis in digital images," *Mathematical Problems in Engineering*, 2016.
- [7] Kumar, S., Mukherjee, S., and Pal, A. K., "An improved reduced feature-based copy-move forgery detection technique," *Multimedia Tools and Applications*, Vol. 82, No. 1, pp. 1431-1456, 2023.
- [8] Alberly, H. A., Hegazy, A. A., and Salama, G. I., "A fast SIFT based method for copy move forgery detection," *Future Computing and Informatics Journal*, Vol. 3, No. 2, pp. 159-165, 2018.
- [9] Wang, C., Zhang, Z., and Zhou, X., "An image copy-move forgery detection scheme based on A-KAZE and SURF features," *Symmetry*, Vol. 10, No. 12, 2018.
- [10] Narayanan, S. S., and Gopakumar, G., "Recursive block based keypoint matching for copy move image forgery detection," In *11th International conference on computing, communication and networking technologies (ICCCNT)*, pp. 1-6, 2020.

- [11] Elaskily, M.A., Elnemr, H.A., Sedik, A., Dessouky, M.M., Banby, G.M.E., Elshakankiry, O.A., and El-Samie, F.E.A., "A novel deep learning framework for copy-moveforgery detection in images," *Multimedia Tools and Applications*, Vol. 79, No. 27, pp. 19167–19192, 2020.
 - [12] Doegar, A., Dutta, M., and Gaurav, K., "Cnn based image forgery detection using pre-trained alexnet model," *International Journal of Computational Intelligence & IoT*, Vol. 2, No. 1, 2019.
 - [13] Albelwi, S., and Mahmood, A., "A framework for designing the architectures of deep convolutional neural networks," *Entropy*, Vol. 19, No. 6, 2017.
 - [14] Chollet, F., "Deep learning with python," Simon and Schuster, 2021.
 - [15] Ramadhani, E., "Photo splicing detection using error level analysis and laplacian-edge detection plugin on GIMP," In *Journal of Physics: Conference Series*, Vol. 1193, No. 1, 2019.
-