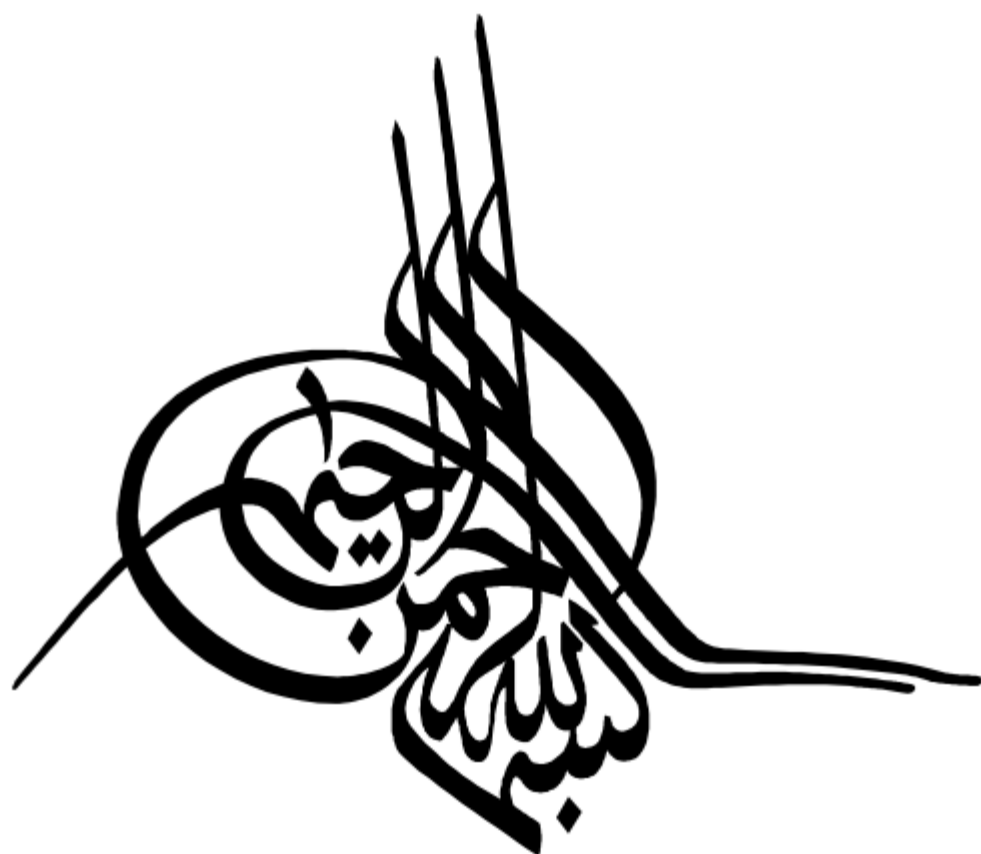






دانشکده: کامپیوتر و فناوری اطلاعات
گروه: هوش مصنوعی

ارائه‌ی روشی مؤثر در تشخیص تصاویر دیجیتالی جعلی

دانشجو: زهرا محمدیان

استاد راهنما:

دکتر علی اکبر پویان

استاد مشاور:

دکتر حمید حسن پور

پایان نامه ارشد جهت اخذ درجه کارشناسی ارشد

شهریور ۱۳۹۲

دانشگاه صنعتی شاهرود

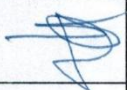
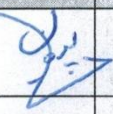
دانشکده: کامپیوتر و فناوری اطلاعات

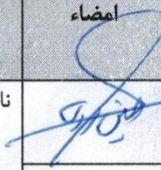
گروه: هوش مصنوعی

پایان نامه کارشناسی ارشد آقای / خانم زهرا محمدیان

تحت عنوان: ارائه‌ی روشی مؤثر در تشخیص تصاویر دیجیتالی جعلی

در تاریخ ۱۳۹۲/۰۶/۲۳ توسط کمیته تخصصی زیر جهت اخذ مدرک کارشناسی ارشد مورد ارزیابی و با درجه عالی مورد پذیرش قرار گرفت.

امضاء	اساتید مشاور	امضاء	اساتید راهنما
	نام و نام خانوادگی: دکتر حمید حسن پور		نام و نام خانوادگی: دکتر علی اکبر پویان
	نام و نام خانوادگی :		نام و نام خانوادگی :

امضاء	نمایندة تحصیلات تکمیلی	امضاء	اساتید داور
	نام و نام خانوادگی : مهندس علی بازقندی		نام و نام خانوادگی : دکتر حسین مروی
			نام و نام خانوادگی : دکتر پیمان کبیری
			نام و نام خانوادگی :
			نام و نام خانوادگی :

این پایان نامه را با نهایت سپاس پیشکش می‌کنم:

به پدرم که با لطف و زحمات بی‌درغش مسیبه شرف و خوشبختی را برایم هموار کرده است و به مادرم که دعای خالصانه اش

همواره بدرقه راه و وجودش دگر می‌من است.

و تقدیم می‌کنم:

به همسر عزیزم که همواره مهربون لطف و محبت خالصانه اش بهتم و خواهم بود.

شکر و قدردانی

حال که به لطف و رحمت لایتنای حضرت حق، مراحل این پایان نامه به اتمام رسیده است، بر خود لازم می دانم تا از همه ی

کسانی که در پیشبرد اهداف این پژوهش مرا کرده اند قدردانی به عمل آورم.

ابتدا از زحمات و پشتیبانی بی دریغ و بی شایه ی استاد ارجمندم، جناب آقای دکتر پویان، که راهنمایی این تحقیق را بر عهده

داشتند کمال شکر را دارم. بی شک بدون حمایت و بهنگری ایشان انجام این پایان نامه مقدور نبود.

و نیز سپاس گذارم از استاد گرامی، جناب آقای دکتر حسن پور، که از هر گونه راهنمایی و مساعدت این جانب در انجام پژوهش
مضايقه نکردند.

در پایان از استاد گران قدر، جناب آقای دکتر زاهدی که سعادت شاکردی ایشان را در دوره ی کارشناسی ارشد داشته ام

قدردانی کرده، از خداوند متعال برای این بزرگوار موفقیت و بهروزی مسألت دارم.

تعهد نامه

اینجانب زهرا محمدیان دانشجوی دوره کارشناسی ارشد رشته هوش مصنوعی دانشکده مهندسی فناوری اطلاعات و کامپیوتر دانشگاه صنعتی شاهرود نویسنده پایان نامه "ارائه‌ی روشی مؤثر در تشخیص تصاویر دیجیتالی جعلی" تحت راهنمایی آقای دکتر علی اکبر پویان متعهد می‌شوم

- تحقیقات در این پایان نامه توسط اینجانب انجام شده‌است و از صحت و اصالت برخوردار است.
- در استفاده از نتایج پژوهش‌های محققان دیگر به مرجع مورد استفاده استناد شده‌است.
- مطالب مندرج در پایان نامه تا کنون توسط خود یا فرد دیگری برای دریافت هیچ نوع مدرک یا امتیازی در هیچ‌جا ارائه نشده‌است.
- کلیه حقوق معنوی این اثر متعلق به دانشگاه صنعتی شاهرود می‌باشد و مقالات مستخرج با نام "دانشگاه صنعتی شاهرود" و یا "Shahrood University of Technology" به چاپ خواهد رسید.
- حقوق معنوی تمام افرادی که در بدست آمدن نتایج اصلی پایان نامه تاثیرگذار بوده‌اند در مقالات مستخرج از پایان نامه رعایت می‌شود.
- در کلیه مراحل انجام این پایان نامه در مواردی که به حوزه اطلاعات شخصی افراد دسترسی یافته یا استفاده شده‌است، ضوابط و اصول اخلاقی رعایت شده‌است.



امضا

تاریخ: ۱۳۹۲/۰۶/۲۳

مالکیت نتایج و حق نشر

کلیه حقوق معنوی این اثر و محصولات آن (مقالات مستخرج، کتاب، برنامه‌های رایانه‌ای، نرم‌افزارها و تجهیزات ساخته‌شده) متعلق به دانشگاه صنعتی شاهرود است. این مطلب باید به نحو مقتضی در تولیدات علمی مربوطه ذکر شود.

استفاده از اطلاعات و نتایج موجود در پایان نامه بدون ذکر مرجع مجاز نمی‌باشد.

چکیده

هدف این پایان نامه، ارائه‌ی روشی برای تشخیص جعل تصاویر از نوع کپی-انتقال^۱ می‌باشد. تشخیص صحت تصاویر، زمانی مهم و حساس می‌شود که بخواهیم از آن تصاویر به عنوان مدرک، در دادگاه‌ها، استفاده کنیم. جعل کپی-انتقال یکی از عمومی‌ترین روش‌های جعل تصاویر می‌باشد که در آن، یک قسمت از تصویر، کپی و در جای دیگری از همان تصویر چسبانده می‌شود.

در این پژوهش به بررسی انواع روش‌های تشخیص جعل کپی-انتقال می‌پردازیم. بسیاری از روش‌ها در صورتی قادر به تشخیص هستند که هیچ تغییری بر روی ناحیه‌ی کپی شده، قبل از چسبانده شدن، صورت نگرفته باشد. اکثر جاعلان برای این‌که ناحیه‌ی کپی شده با نواحی اطرافش همخوانی داشته باشد و تصویر طبیعی‌تر به نظر برسد، بر روی ناحیه‌ی کپی شده، یکسری تبدیلات هندسی اعمال می‌کنند.

یک روش تشخیص جعل خوب باید نسبت به انواع تغییرات و تبدیلات هندسی مقاوم باشد. بسیاری از روش‌ها نسبت به تمام یا برخی از تغییرات مقاوم نیستند یا از لحاظ محاسباتی بسیار پیچیده و زمانبر هستند. روش تشخیص جعل با استفاده از تبدیل ویژگی‌های مستقل از مقیاس^۲ نسبت به انواع تغییرات مقاوم است و زمان پردازش پایینی دارد، اما این روش در صورتی که جعل در ناحیه‌ی هموار تصویر صورت گرفته شده باشد، قادر به تشخیص نخواهد بود. این اشکال را توسط ممان‌های زرنیک برطرف کرده‌ایم. در این تحقیق، به ارائه‌ی روشی ترکیبی برای تشخیص جعل کپی-انتقال خواهیم پرداخت که از ترکیب دو روش تشخیص جعل بر پایه‌ی ممان‌های زرنیک و ویژگی‌های SIFT حاصل شده است. باتوجه به خصوصیات این دو روش، روش پیشنهادی در برابر انواع تبدیلات و تغییرات هندسی مقاوم و قادر به تشخیص جعل در نواحی هموار می‌باشد همچنین بهبودهای زیر بر روی آن‌ها اعمال شده‌است:

✓ تشخیص چندین ناحیه‌ی کپی شده از یک ناحیه. در حالی‌که روش‌های قبلی، قادر به تشخیص فقط یک ناحیه‌ی کپی شده از ناحیه‌ی اصلی هستند.

✓ این روش بین اشیا و نواحی شبیه به هم که منتج به استخراج ویژگی‌های مشابه می‌شود، و نواحی یا اشیا کپی شده تمییز قائل می‌شود.

✓ همچنین در صورت تشخیص جعل، به تعیین پارامترهای تبدیلات هندسی و در انتها به مقایسه‌ی روش پیشنهادی با سایر روش‌ها خواهیم پرداخت.

کلمات کلیدی: تشخیص جعل، کپی-انتقال، تصاویر دیجیتالی، ویژگی SIFT، ممان زرنیک

^۱ Copy-move

^۲ Scale Invariant Features Transform (SIFT)

مقالات استخراج شده از پایان نامه

1. Pouyan A. A. and Mohamadian Z. (2012) "Image duplication forgery detection using SIFT features and Zernik moments" **WULFENIA**, 19, 10, pp **103-115**.
2. Pouyan A.A. and Mohamadian Z. (2013) "Detection of Duplication Forgery in Digital Images in Uniform andNon-Uniform Regions" **UKSim-AMSS 15th International Conference on Modelling and Simulation, Cambridge University, Emmanuel College**.

فهرست مطالب

صفحه	عنوان
۳	فصل اول: مقدمه
۴	۱-۱- اهمیت تشخیص تصاویر جعلی
۵	۱-۲- انواع جعل داده‌های چند رسانه‌ای
۶	۱-۳- چند مثال واقعی از جعل تصاویر
۸	فصل دوم: روش‌های جعل و روش‌های تشخیص جعل
۸	۲-۱- نمونه‌هایی از روش‌های جعل و دستکاری در تصاویر دیجیتال
۱۱	۲-۲- جعل محتوای ویدئویی
۱۱	۲-۳- روش‌های تشخیص تصاویر جعلی
۱۲	۲-۳-۱- روش‌های بر پایه‌ی واتر مارکینگ دیجیتال
۱۴	۲-۳-۲- روش‌های بر پایه‌ی مقدار درهم
۱۵	۲-۳-۳- روش‌های بر پایه‌ی مالتی مدیای قانونی
۱۶	۲-۳-۳-۱- تشخیص بر پایه‌ی همبستگی
۱۸	۲-۳-۳-۲- فشرده سازی JPEG مضاعف
۲۰	۲-۳-۳-۳- تشخیص بر پایه‌ی ناهماهنگی در نور
۲۱	۲-۳-۳-۴- تشخیص بر پایه‌ی تیزی و مات شدگی تصویر
۸	فصل سوم: جعل کپی - انتقال
۲۵	۳-۱- فرآیند تشخیص جعل کپی - انتقال
۲۷	۳-۲- روش‌های بر پایه‌ی بلوک بندی و روش‌های برپایه‌ی نقاط کلیدی
۲۸	۳-۲-۱- الگوریتم‌های بر مبنای بلوک بندی
۳۰	۳-۲-۲- الگوریتم‌های بر مبنای نقاط کلیدی
۲۵	فصل چهارم: بررسی کارهای انجام شده در زمینه‌ی تشخیص جعل کپی - انتقال
۳۲	۴-۱- بررسی مقالات و کارهای انجام شده
۳۳	۴-۲- ویژگی‌ها و مقایسه‌ی آن‌ها در جعل کپی - انتقال

فصل پنجم: ویژگی‌های SIFT و ممان‌های زرنیک	۳۲
۱-۵- ویژگی‌های مقاوم در برابر تغییر سایز (SIFT)	۳۸
۱-۱-۵- تعیین نقاط مینیمم و ماکزیمم در فضای مقیاس	۳۸
۲-۱-۵- انتخاب نقاط کلیدی از بین نقاط ماکزیمم و مینیمم اصلی	۴۰
۳-۱-۵- تعیین موقعیت مکانی هر نقطه‌ی کلیدی	۴۱
۴-۱-۵- تشکیل توصیف کننده برای هر نقطه‌ی کلیدی	۴۱
۲-۵- ممان‌های زرنیک	۴۲
۱-۲-۵- تعریف	۴۳
۲-۲-۵- مقاوم بودن ممان‌های زرنیک در برابر چرخش	۴۴
فصل ششم: روش پیشنهادی و مقایسه‌ی آن با روش‌های دیگر	۴۶
۱-۶- تشخیص جعل کپی- انتقال با استفاده از ویژگی‌های SIFT	۴۸
۱-۱-۶- انطباق نقاط کلیدی	۴۸
۲-۱-۶- خوشه بندی و تشخیص جعل	۵۰
۳-۱-۶- تخمین تبدیلات هندسی	۵۳
۲-۶- پایگاه داده	۵۵
۳-۶- کشف جعل	۵۶
۴-۶- نتایج به دست آمده	۵۹
۵-۶- مقایسه‌ی روش پیشنهادی با سایر روش‌ها	۵۹
۶-۶- تشخیص جعل کپی- انتقال با استفاده از ممان‌های زرنیک	۷۰
۷-۶- نتیجه‌ی حاصل از ترکیب ویژگی‌های SIFT و ممان‌های زرنیک	۷۲
۸-۶- بررسی میزان مقاومت روش پیشنهادی در برابر فشرده سازی JPEG، افزودن نویز و اصلاح گاما	۷۳
فصل هفتم: نتیجه گیری و پیشنهادها برای پژوهش‌های آینده	۴۸
پژوهش‌های آینده	۴۸
۱-۷- نتیجه گیری	۷۴
۲-۷- پیشنهادها برای پژوهش‌های آینده	۷۶
مراجع	۷۷

فهرست شکل ها

صفحه	عنوان
۵.....	شکل (۱-۱) : نمونه‌ای از جعل تصاویر تاریخی
۹.....	شکل (۱-۲): نمونه‌ای از حذف
۱۰.....	شکل (۲-۲): نمونه‌ای از جایگزینی
۱۰.....	شکل (۳-۲): نمونه‌ای از تکرار
۱۰.....	شکل (۴-۲): نمونه‌ای از مونتاژ
۱۱.....	شکل (۵-۲): نمونه‌ای از تصویر تولیدشده توسط گرافیک کامپیوتری.
۱۲.....	شکل (۶-۲): روش تشخیص جعل بر پایه‌ی واتر مارکینگ
۱۴.....	شکل (۷-۲): روش تشخیص جعل بر پایه‌ی مقدار درهم.....
۱۵.....	شکل (۸-۲): روش تشخیص جعل بر پایه‌ی مالتی مدیای قانونی.
۱۷.....	شکل (۹-۲): تشخیص بازنمونه برداری
۱۸.....	شکل (۱۰-۲): تشخیص جعل براساس روش درونیابی آرایه‌ی فیلتر رنگی
۱۹.....	شکل (۱۱-۲): تشخیص جعل با توجه به فشرده سازی JPEG مضاعف.
۲۱.....	شکل (۱۲-۲): تشخیص جعل براساس ناهم‌هنگی نوری.....
۲۲.....	شکل (۱۳-۲): نمونه‌ای از تشخیص جعل برپایه‌ی تیزی و مات شدگی تصویر
۲۴.....	شکل (۱-۳): طرح کلی روش‌های فعال و غیر فعال.
۲۵.....	شکل (۲-۳): نمونه‌ای از جعل کپی- انتقال برای پوشاندن یک شیء.
۲۵.....	شکل (۳-۳): فرآیند مشترک برای تشخیص جعل کپی- انتقال.
۴۰.....	شکل (۱-۵): هرم تصویری با تابع گوسین
۴۲.....	شکل (۲-۵): توصیف کننده‌های محلی در تصویر برای یک نقطه‌ی کلیدی
۵۹.....	شکل (۱-۶): نمونه‌ای از تشخیص جعل توسط روش پیشنهادی
۵۹.....	شکل (۲-۶): تغییر سایز در ناحیه‌ی جعل شده.....
۶۰.....	شکل (۳-۶): چرخش در ناحیه‌ی جعل شده.....
۶۱.....	شکل (۴-۶): چند ناحیه‌ی کپی شده از یک ناحیه.
۶۲.....	شکل (۵-۶): افزودن نویز گوسین سفید.
۶۳.....	شکل (۶-۶): فشرده سازی JPEG تصویر
۶۴.....	شکل (۷-۶): مات کردن تصویر
۶۵.....	شکل (۸-۶): بهبود توسط اصلاح گاما.....
۶۶.....	شکل (۹-۶): تبدیل افاین در ناحیه‌ی جعل شده.....
۶۷.....	شکل (۱۰-۶): تبدیل افاین در ناحیه‌ی جعل شده

- شکل (۶-۱۱): نواحی شبیه به هم..... ۶۷
- شکل (۶-۱۲): استخراج ویژگی‌های SIFT..... ۶۸
- شکل (۶-۱۳): تشخیص جعل با استفاده از فقط ویژگی‌های SIFT..... ۶۸
- شکل (۶-۱۴): تشخیص نواحی هموار با استفاده از روش پیشنهادی..... ۷۲

فهرست جداول

عنوان	صفحه
جدول (۱-۳): گروه‌بندی مجموعه ویژگی‌ها برای تشخیص جعل کپی- انتقال	۲۸
جدول (۱-۶): ۱۰ ترکیب مختلف از حملات اعمال شده بر روی ناحیه‌ی اصلی بر روی تصاویر پایگاه داده	۵۶
جدول (۲-۶): فاز آموزش: مقادیر TPR و FPR با توجه به T	۵۷
جدول (۳-۶): خطای تخمین پارامترهای هندسی بر روی تصاویر پایگاه داده‌ی MICC-F220	۵۸
جدول (۴-۶): ارزیابی عملکرد تشخیص در برابر اصلاح گاما، فشرده سازی JPEG و افزودن نویز اعمال شده بر روی ناحیه‌ی جعل شده و تبدیل یافته	۷۳

فصل اول

مقدمه

امروزه به دلیل در دسترس بودن نرم افزارهای ویرایش و پردازش تصویر، می‌توان بدون به‌جا گذاشتن هرگونه ردپایی، ویژگی‌های مهمی به تصویر اضافه یا کم کرد و تصاویر را جعل نمود [۱]. به دلیل پیشرفت ابزارهای تولید تصاویر جعلی، عملیات جعل تصاویر از کنترل متخصصان خارج شده [۲] و حتی افرادی که در این زمینه متخصص نیستند می‌توانند به راحتی، تصاویر را تغییر دهند [۳]. این موضوع مزایای زیادی دارد، ولی از طرفی ضررهای بیشتری را به دنبال خواهد داشت که در زیر به برخی از آن‌ها اشاره خواهیم کرد.

۱-۱- اهمیت تشخیص تصاویر جعلی

به دلیل وجود عکس‌های دیجیتالی در همه جا مانند جلد مجله‌ها، روزنامه‌ها، دادگاه‌ها و سرتاسر اینترنت، ما در تمام لحظات با آن‌ها روبرو می‌شویم و به آنچه که می‌بینیم اعتماد می‌کنیم [۱]. در دسترس بودن و عام پسند بودن تکنولوژی تولید تصاویر دیجیتال (مانند: دوربین‌ها، کامپیوترها و نرم‌افزارها) و استفاده وسیع آن‌ها در اینترنت، عکس‌های دیجیتالی را به منابع اطلاعات تبدیل کرده است و با پیشرفت تکنولوژی دست‌کاری تصاویر دیجیتالی، تصاویر دیجیتالی فریب دهنده، رشد سریعی پیدا کرده‌اند [۴].

در سال‌های اخیر، تصاویر جعلی، روی علم، قانون، سیاست، رسانه‌ها و تجارت، تاثیر بسزایی گذاشته است [۱].

در سال ۲۰۰۴ م، ۵ میلیون کاربر عکس‌های خود را که توسط نرم افزارهای موجود، مونتاژ کرده بودند، در سایت مربوط به مسابقه^۱ برای کسب جوایز برترین عکس، به نمایش گذاشتند. کیفیت بالای برخی از عکس‌ها تشخیص اصلی یا جعلی بودن آن‌ها را دشوار کرده بود [۵]، [۶] و [۷]. در این

^۱ www.worth1000.com

مسابقه، بعضی از افراد، عکس‌های جعلی را در سایت قرار دادند و باعث شدند که مسابقه از حالت منصفانه بودن خارج شود.

ممکن است اشخاصی چهره‌ی یک فرد را به جای فرد دیگری قرار دهند و تصویرشان را برای خدشه دار کردن شهرت و آبروی آنان، در اینترنت و در دسترس عموم بگذارند. بنابراین، با توجه به آسان شدن روش‌های دست‌کاری عکس‌ها، نباید به سادگی به هر چیزی که می‌بینیم اعتماد کنیم.

با توجه به مشکلات مطرح شده، مشخص است که یکی از مسائل اساسی در بحث قانونی تصاویر^۱، مشخص کردن اصلی یا جعلی بودن تصاویر است. موضوع زمانی جدی می‌شود که بخواهیم از این تصاویر به عنوان مدرک، در دادگاه‌ها استفاده کنیم. بنابراین، برای حفظ حقوق شخصی افراد، ضروری است که به روش‌هایی که می‌توانند موارد جعلی را تشخیص دهند دسترسی داشته باشیم [۸].

در هر صورت، نیاز جدی به ابزارهای تشخیص جعل وجود دارد تا بار دیگر اعتماد مردم را به تصاویر منتشر شده جلب کنیم. این ابزارها و روش‌ها در مقالات زیادی ارائه شده است که در این پایان‌نامه به آن‌ها نیز خواهیم پرداخت.

۲-۱- انواع جعل داده‌های چند رسانه‌ای

داده‌های چند رسانه‌ای (مالتی مدیا)^۲ شامل عکس، صوت و ویدئو توسط یکی از عملیات زیر

جعل می‌شوند:

✓ حذف^۳؛

✓ جایگزینی^۴؛

✓ تکرار^۵؛

^۱ image forensics

^۲ multimedia data

^۳ data removal

^۴ replacement

^۵ Replication or Copy- move

✓ مونتاز^۱؛

✓ تصاویر گرافیکی مصنوعی تولید شده توسط کامپیوتر^۲.

که در فصل دوم به اختصار به هر کدام خواهیم پرداخت.

روش‌های تشخیص جعل موجود به سه گروه تقسیم می‌شوند:

روش‌های بر پایه واتر مارکینگ^۳، روش‌های بر پایه مقدار درهم^۴، روش‌های بر پایه مالتی مدیای قانونی^۵.

هر کدام از این روش‌ها در سطوح مختلفی از دقت و بازدهی به کار گرفته می‌شوند. به این روش‌ها نیز خواهیم پرداخت.

۳-۱- چند مثال واقعی از جعل تصاویر

فن ساختن عکس‌های جعلی به قدمت عکاسی است. در همان سال‌های آغازین اختراع دوربین عکاسی، عکس گرفتن روش انتخابی مردم برای ساختن پورتره^۶ شد و عکاسان از همان ابتدا دریافته بودند که با دست‌بردن در عکس‌های خود و راضی کردن مشتری، می‌توانند میزان فروش خود را افزایش دهند.

در جنگ داخلی ایالات متحده آمریکا^۷، برای تاثیر گذاری بیشتر، به بسیاری از عکس‌ها جزئیاتی را اضافه کردند. اولین عکاسان، روش ترکیب کردن عکس‌ها را نیز امتحان کردند. ابتدایی ترین نمونه از این روش را در شکل (۱-۱) مشاهده می‌کنید [۹]. در این تصویر، ژنرال فرانسیک در سمت راست ترین نقطه عکس، اضافه شده در حالی که او در تصویر اصلی وجود ندارد. مثال‌های بسیار زیادی از جعل تصاویر در سال‌های اولیه عکاسی وجود دارد و هدف تمام آن‌ها بهبود کیفیت عکس و

¹ photomontage

² computer-aided media generation

³ watermarking-based scheme

⁴ perceptual hash-based

⁵ multimedia forensic-based scheme

⁶ portraits

⁷ Civil War

یا اضافه کردن مایه‌های طنز و شوخی به عکس بوده است و نه فریب دادن دیگران، اگرچه در اواسط قرن بیستم عکاسان دریافتند که جعل تصاویر می‌تواند ابزار قدرتمندی برای تغییر دادن درک عمومی و یا حتی تاریخ باشد.

ارتش نازی آلمان به خاطر تبلیغاتش معروف است و عکس‌های دست‌کاری شده زیادی را برای فریب دادن مردم منتشر کرده‌است. در ردیف پایین شکل (۱-۱) جوزف گوبل، یکی از وزرای هیتلر از تصویر اصلی حذف شده است [۱۰].



(ب)



(الف)



(د)



(ج)

شکل (۱-۱): نمونه‌ای از جعل تصاویر تاریخی. الف) تصویر اصلی، ب) تصویر جعلی (ژنرال فرانسیک، به تصویر اصلی اضافه شده است)، ج) تصویر اصلی، د) تصویر جعلی (ژنرال جوزف گوبل از تصویر اصلی حذف شده است).

مثال‌های مشابهی از کشورهای آمریکا و روسیه وجود دارد که در آنجا تصاویر افراد با چهره‌ی نازیبا را حذف می‌کنند و یا این‌که تصاویر افرادی را به دلایل سیاسی، به عکس‌ها اضافه می‌کنند؛ ولی با وجود مثال‌های بیشماری از تصاویر جعلی، بسیاری از افراد بر این باورند که "تصاویر دوربین هرگز به آنها دروغ نمی‌گوید."

در سال ۲۰۰۴ م.، یک تیم پزشکی به رهبری یک دانشمند کره‌ای به نام دکتر هوانگ^۱، نتایج اولیه تحقیقات در مورد سلول‌های استم^۲ را در مجله علمی به چاپ رساندند، ولی بعدها مشخص شد که عکس‌های منتشر شده در مقاله‌ی آنان، جعلی بوده است [۱۱]، [۱۲] و [۱۳] با مشخص شدن این ماجرا، این مقاله از تمام مجلات جمع آوری شد و دکتر هوانگ از سمت خود استفا داد. از آن پس، داوران این مجله^۳ اصل یا جعل بودن تمام عکس‌های مقالات دریافتی از سال ۲۰۰۲ به بعد را مورد بررسی قرار دادند و دریافتند که حدود ۲۵ درصد آن‌ها جعلی بوده است [۱۴].

مثالی دیگر از جعل دیجیتالی توسط دکتر توماس پوگیو^۴ در یک کنفرانس الکترونیک^۵ در سال ۲۰۰۳ م. در سانتا کلارا^۶ ارائه شد. در این نشست، دکتر پوگیو نشان داد مهندسين می‌توانند حرکات لب انسان را در یک کلیپ کوتاه ویدئویی بیاموزند و این حرکات لب را دست‌کاری کنند تا بتوانند یک محتوای صوتی دلخواه را جایگزین محتوای قبلی کنند. به عنوان یک مثال خوب، یک کلیپ ویدئویی تبلیغاتی نشان داده شد که حرکات لب‌های خواننده را طوری تغییر دادند که گویی یک شعر مشهور را می‌خواند و آن شعر را جایگزین شعر تبلیغاتی کردند!

¹ Dr. Hwang Woo-Suk

² stem cells

³ Journal of Cell Biology

⁴ Dr. Tomaso A. Poggio

⁵ Electronic Imaging

⁶ Santa Clara

فصل دوم

روش‌های جعل و روش‌های تشخیص جعل

همان‌گونه که گفته شد، در عصر دیجیتال که روش‌های مالتی مدیا به سرعت رشد می‌کنند، داده‌های چند رسانه‌ای مانند: عکس، صوت و ویدئو به آسانی می‌توانند توسط نرم افزارهای رایج، جعل شوند. چند نمونه از این نرم افزارها عبارتند از: فتوشاپ^۱، WaveCN یا FFmpeg [۵].

بنابراین، ایجاد تصاویر ساختگی برای هرکسی امکان پذیر شده است. اگر تصویری که به عنوان یک مدرک استفاده می‌شود، دستکاری شود آنگاه در بسیاری از کاربردها قابل استفاده نخواهد بود. از این رو نیاز مبرم به بررسی معتبر بودن تصاویر دیجیتال با عدم دسترسی به تصویر اصلی وجود دارد. برای این کار ابزارها و روش‌هایی مطرح شدند که به بررسی صحت و یا عدم صحت تصاویر می‌پردازند. در ادامه به بررسی انواع روش‌های جعل و روش‌های تشخیص آن خواهیم پرداخت.

۲-۱- نمونه‌هایی از روش‌های جعل و دست‌کاری در تصاویر دیجیتال

به طور کلی روش‌هایی که محتوا را تغییر می‌دهند^۲ به گروه‌های زیر تقسیم می‌شوند [۱۸]:

حذف کردن^۳: این روش شامل عملیاتی همچون حذف کردن قسمتی از محتوای مالتی مدیا می‌باشد. حذف یک ماشین در حالت حرکت از یک تصویر، بریدن یک قسمت از یک دنباله صوت و حذف یک شخص از یک فریم ویدئویی، نمونه‌هایی از این روش هستند. عموماً برای به‌دست آوردن کیفیت بهتر یا تاثیر گذاشتن روی محتوا، این روش با روش‌های دیگری مانند فیلتر کردن و حذف نویز ترکیب می‌شود.

جایگزینی: این گروه شامل عملیاتی است که بخش‌هایی از محتوای مالتی مدیا را با بخش-هایی از محتوای دیگر جایگزین می‌کند. جایگزین کردن صورت یک شخص در یک عکس با صورت شخصی دیگر در عکسی دیگر، جایگزین کردن یک قسمت از یک داده‌ی صوتی با قسمت دیگری از

¹ Photoshop

² content-altering operations

³ Removing

صوتی دیگر، مثال‌هایی از این روش هستند. این جایگزینی‌ها عموماً با ترکیب چندین روش مانند بریدن^۱، چسباندن^۲ و هموار سازی^۳ به دست می‌آید.

تکرار: این گروه شامل عملیاتی است که تعداد اشیا را در یک محتوا با استفاده از عملیات کپی کردن و چسباندن آن‌ها از یک مکان به مکانی دیگر در همان محتوا افزایش می‌دهند. به طور کلی این روش با روش‌های دیگری مانند کپی کردن، چسباندن، و صاف کردن ترکیب می‌شوند. در این پایان نامه به بررسی روش‌های تشخیص جعل تکرار یا کپی - انتقال یا کپی - چسباندن می‌پردازیم.

مونتاژ: این گروه شامل روش‌هایی است که تصاویر مختلفی را ترکیب کرده و در نهایت یک تصویر جدید با کیفیت بالا تولید می‌کند [۱۵].

تصاویر گرافیکی مصنوعی تولید شده توسط کامپیوتر: این گروه شامل روش‌هایی است که تصاویر گرافیکی مصنوعی را توسط نرم افزارهای گرافیکی تولید می‌کنند. این تصاویر از روی نمونه واقعی آن‌ها شبیه سازی می‌شوند. کارتون سازی^۴ نمونه ای از این روش است که تصاویر طبیعی دیجیتال یا ویدئویی را به کارتون تبدیل می‌کند [۱۶]. به طور کلی کارتون سازی محتوای مدیا با محتوای اصلی‌اش متفاوت است.

در ادامه چند نمونه از جعل تصاویر آورده شده است.



تصویر جعلی



تصویر اصلی

شکل (۱-۲): نمونه‌ای از حذف. در شکل سمت راست تصویر شخص حذف شده است [۱۷]

¹ wiping

² pasting

³ Smoothing

⁴ Cartoonization



تصویر جعلی



تصویر اصلی

شکل (۲-۲): نمونه‌ای از جایگزینی. در تصویر جعلی، صورت شخصی جایگزین صورت فرد تصویر اصلی شده است [۱۷].



تصویر جعلی



تصویر اصلی

شکل (۳-۲): نمونه‌ای از تکرار. تصویر جعلی حاصل از کپی کردن یک هواپیما و چسباندن آن در همان عکس [۳].

در سال ۲۰۰۶م. تصویر شکل (۲-۴) که شامل آهوها در جلوی عکس و تصویر قطار در پس زمینه است، برنده زیباترین عکس در CCTV^۱ شد. اما بعدها مشخص شد که تصویر آهوها، قطار و حتی سنگ ریزه‌ها از عکس‌های دیگر مونتاژ شده است.



شکل (۴-۲): نمونه‌ای از مونتاژ. تصویر آهوها، قطار و حتی سنگ ریزه‌ها از عکس‌های دیگر چسبانده شده‌اند [۵].

¹ Chinese Central Television Station

تصویر شکل (۵-۲)، چهره‌ی یکی از هنر پیشه‌های معروف کره‌ای است که توسط یک گرافیکست ماهر توسط نرم افزارهای گرافیکی کامپیوتری کشیده شده‌است.



شکل (۵-۲): نمونه‌ای از تصویر تولید شده توسط گرافیک کامپیوتری [۵].

۲-۲- جعل محتوای ویدئویی

برای دست‌کاری کردن یک ویدئوی آنالوگ می‌توانیم به آسانی آن را دیجیتالی^۱ کنیم، برای این کار ابتدا داده را وارد کامپیوتر کرده و عملیات جعل را روی آن انجام می‌دهیم و در مرحله آخر، نتایج به‌دست آمده را به فرمت NTSC در می‌آوریم.

۲-۳- روش‌های تشخیص تصاویر جعلی

روش‌های تشخیص جعل، عملیات مخربی را که محتوای داده‌های مالتی مدیا را دست‌کاری می‌کنند و یا کپی‌های جعلی می‌سازند شناسایی می‌کنند. همان‌طور که قبلاً ذکر شد، این روش‌ها عبارتند از: روش‌های بر پایه‌ی واتر مارکینگ دیجیتال [۱۸] و [۱۹]، روش‌های بر پایه‌ی درهم سازی [۲۰] و [۲۱] و روش‌های بر پایه‌ی مالتی مدیای قانونی.

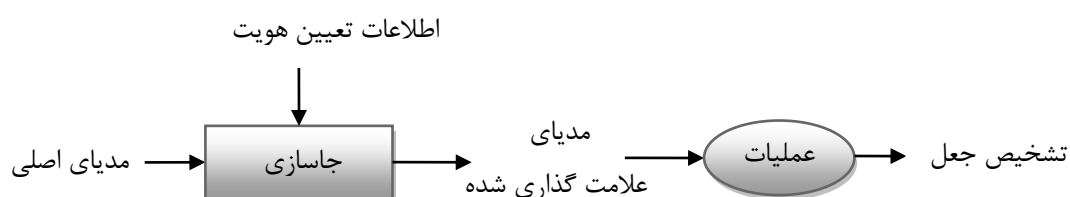
¹ digitize

۲-۳-۱- روش‌های بر پایه‌ی واتر مارکینگ دیجیتال

واترمارکینگ دیجیتال اطلاعات صحیح و معتبری را در هنگام تولید محتوای داده، درون آن جاسازی می‌کند [۵]. به عنوان مثال، فرآیند جاسازی اطلاعات واترمارک، یکی از ویژگی‌هایی است که در دوربین دیجیتال ساخته می‌شود. همان‌طور که در شکل (۲-۶) نشان داده شده‌است، در این روش اطلاعات واتر مارک مانند درستی و صحت هویت مالک اصلی تصویر و... به طور غیر قابل مشاهده، در محتوای مالی مدیا با دست‌کاری مستقیم محتوای آن جاسازی شده‌است. این عملیات جاسازی، به هنگام تولید محتوای مالی مدیا انجام می‌شود (مانند دوربین‌ها) [۱۸]. اطلاعات جاسازی شده از محتوای مالی مدیا استخراج شده و با محتوای اصلی مقایسه می‌شود. این مقایسه مشخص می‌کند که آیا تصویر جعل شده است یا خیر و حتی ناحیه جعل شده را جدا می‌کند.

این روش دارای دو عیب آشکار است:

۱. عملیات جاسازی اطلاعات در اغلب اوقات در کاربردهای عملی، غیر قابل دسترسی است.
۲. عملیات جاسازی اطلاعات، کیفیت محتوای مالی مدیا را کاهش می‌دهد که این امر برای بسیاری از کاربردها غیر قابل قبول است.



شکل (۲-۶): روش تشخیص جعل بر پایه واتر مارکینگ

در حالی که واترمارکینگ دیجیتال، می‌تواند اطلاعات مفیدی در مورد درستی تصویر و تاریخچه-ی پردازش آن در اختیار ما بگذارد، ولی همان‌طور که گفته شد، واترمارک باید قبل از این که جعل و

دست‌کاری در تصویر صورت بگیرد، در آن جاسازی شود با وجود این که تمام وسایل دیجیتالی، مجهز به یک تراشه واترمارکینگ هستند، کشف تمام جعل‌های موجود با استفاده از واترمارک بعید است.

واترمارکینگ نیاز به دوربین‌های مخصوصی دارد مانند: ^۱ Canon EOS-1D یا Nikon D2Xs [۲].

هر دو دوربین یک توضیح مختصر منحصر به فرد^۲ برای هر تصویر تولید کرده و آن را به تاریخ گرفته شدن عکس مجهز می‌کنند.

برای تشخیص صحت یا عدم صحت تصویر در تاریخی دیگر، توضیح منحصر به فرد را دوباره تولید و با مقدار اصلی مقایسه می‌کنیم. در صورت تفاوت بین این دو مقدار، جعل تشخیص داده می‌شود. در حالی که این دوربین‌ها می‌توانند در بعضی موارد مانند کاربردهای قانونی قابل استفاده و مفید باشند، محدودیت‌های قابل توجهی دارند. مهمترین محدودیت این است که امروزه دوربین‌های بسیار کم و گران قیمت، این ویژگی را دارند و در آینده، این سیستم‌ها اجازه نمی‌دهند که تغییراتی شامل بهبود کیفیت عکس مانند شارپ کردن^۳ و یا افزایش کنتراست^۴ روی تصاویر اعمال شود.

روش‌های بسیار دیگری برای استفاده از واترمارکینگ وجود دارد، بعضی از آن‌ها اجازه تغییرات را می‌دهند و بقیه برای آشکار سازی تغییرات انجام شده طراحی شده‌اند [۲۲]، [۲۳] و [۲۴]. به عنوان مثال، واتر مارک‌های نیمه ضعیف^۵، به تغییرات ساده‌ای مانند فشردن سازی JPEG اجازه می‌دهند که بر روی تصویر انجام شوند، در حالی که واترمارک‌های تل-تایل^۶ می‌توانند برای تشخیص دست‌کاری تصویر، تحلیل و مورد بررسی قرار گیرند. تمام روش‌های واترمارکینگ نیاز به یک سخت افزار یا نرم افزار مخصوص برای جاسازی واترمارک در تصاویر دارند و بعید است که تولید کنندگان با قرار دادن این تکنولوژی در تمام دوربین‌هایی که تولید می‌کنند موافق باشند. بنابراین، واترمارکینگ دیجیتال محدود به دامنه‌ی مسئله است، جایی که ساخت و مدل دوربین‌ها بتوانند کنترل شوند.

^۱ EOS-1D Mark II

^۲ image-specific digest

^۳ sharpening

^۴ Contrast enhancing

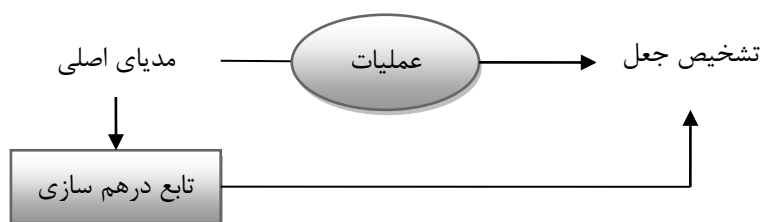
^۵ semi-fragile watermarks

^۶ tell-tale watermarks

۲-۳-۲- روش‌های بر پایه‌ی مقدار درهم

همان‌طور که در شکل (۷-۲) نشان داده شده، تابع درهم سازی بر روی محتوای مالتی مدیا عمل می‌کند [۵]. یک مقدار درهم، شامل رشته‌ای با یک طول مشخص است که توسط شخص تعیین هویت کننده نگهداری می‌شود. مقدار درهم جدید از محتوای مالتی مدیا استخراج شده و با مقدار ذخیره شده مقایسه می‌شود و اگر تفاوتی مشاهده شد، جعل صورت گرفته است. مانند روش واترمارکینگ، در این روش نیز مقدار درهم در هنگام تولید محتوای مالتی مدیا تولید می‌شود (مانند دوربین‌ها). تفاوت این دو روش در این است که مقدار درهم، محتوای مالتی مدیا را تغییر نمی‌دهد، اما اشکال این روش نیز مانند روش واترمارکینگ این است که عملیات محاسبه مقدار درهم که در حین ایجاد مدیا انجام می‌شود، در کاربردهای عملی، غیر قابل دسترسی است.

همان‌طور که گفته شد، در این روش، مقدار درهم^۱ از محتوای مالتی مدیا و در حین تولید محتوای مالتی مدیا، محاسبه می‌شود. بنابراین، هم واترمارک و هم مقدار درهم باید در محتوای مالتی مدیای اصلی جاسازی شوند. در عمل ممکن است با این کار کیفیت مالتی مدیا تنزل پیدا کند و یا ابزارهایی که برای این کار مناسبند در دسترس نباشند.



شکل (۷-۲): روش تشخیص جعل بر پایه‌ی مقدار درهم

^۱ hash value

۲-۳-۳- روش‌های بر پایه‌ی مالتی مدیای قانونی

مالتی مدیای قانونی^۱ [۱۷] و [۱۸]، اطلاعات ارزشمندی را از محتوای داده‌ی مالتی مدیا استخراج و از آن‌ها برای تایید صحت محتوا استفاده می‌کند. روش‌های مالتی مدیای قانونی شامل تشخیص منبع رسانه (موبایل، دوربین دیجیتال، اسکنر و...)، تشخیص جعل و... هستند [۵]. در این روش، تشخیص جعل با استفاده از ویژگی‌های مهم تصاویر صورت می‌گیرد که اغلب با استفاده از تست‌ها یا آموزش‌های آماری به‌دست می‌آیند، سپس با استفاده از این ویژگی‌ها، به تشخیص اشیای غیر معمول و یا نواحی دست‌کاری شده می‌پردازند. عملکرد مالتی مدیای قانونی و میزان دقت تشخیص، بستگی به ویژگی‌های انتخابی دارد. بهترین روش انتخابی، منجر به بیشترین میزان دقت تشخیص می‌شود.

با توجه به موارد گفته شده، در شرایط بدون انجام پیش پردازش، فقط روش‌های قانونی می‌توانند جعل را تشخیص دهند. این روش‌ها ویژگی‌هایی را که می‌توانند تصویر اصلی را از تصویر دست‌کاری شده تمایز دهند، استخراج می‌کنند. بعضی از این روش‌ها با استفاده از ویژگی‌هایی مانند: ویژگی همبستگی، فشردگی، فشرده سازی مضاعف، ویژگی‌های نوری و ویژگی‌های آماری داده (مانند میزان تیزی و مات شدگی^۲ تصویر) و برخی دیگر با استفاده از توابعی مانند (تابع تشخیص تکرار، تشخیص مونتاژ و تشخیص عکس‌های مصنوعی) به تشخیص عکس‌های اصلی از جعلی می‌پردازند که در بخش بعدی، به برخی از آن‌ها اشاره خواهیم کرد. شمای کلی این روش‌ها در شکل (۲-۸) آورده شده است [۵].



شکل (۲-۸): روش تشخیص جعل بر پایه‌ی مالتی مدیای قانونی.

به تازگی، استفاده از این روش‌ها، که مناسب‌تر از روش‌های واترمارکینگ و درهم سازی است، محققان زیادی را علاقه‌مند کرده است.

^۱ multimedia forensics

^۲ Blurriness

۲-۳-۳-۱- تشخیص برپایه‌ی همبستگی^۱

برخی همبستگی‌ها که بین پیکسل‌های همسایه در حوزه‌ی مکان^۲ و زمان وجود دارد اغلب درحین تولید محتوای مالتی مدیا تولید می‌شود. از نقطه نظر محتوای مالتی مدیا، این همبستگی‌ها قابل شناسایی و استفاده در تشخیص جعل می‌باشند. به‌طور کلی، دو نوع از عملیات اغلب مورد توجه قرار می‌گیرند، بازنمونه‌برداری^۳ و درونیابی آرایه فیلتر رنگی (CFA)^۴.

۲-۳-۳-۱-۱- تشخیص بازنمونه‌برداری

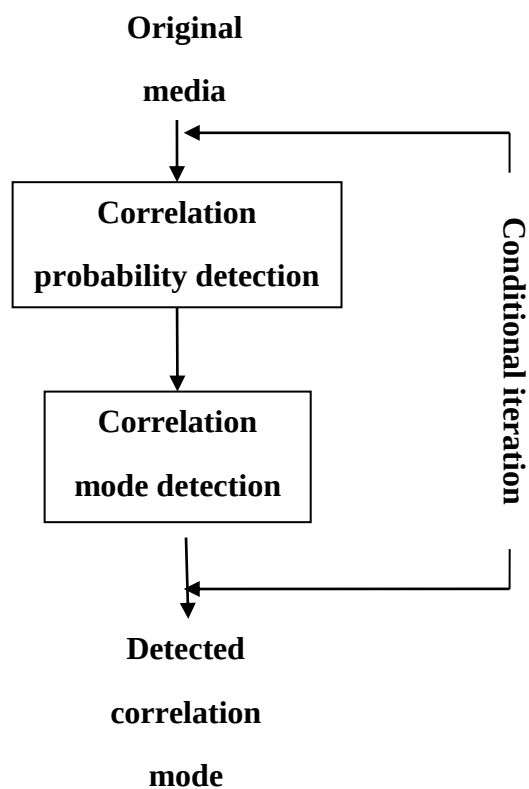
بازنمونه‌برداری اغلب در خلال عملیات‌های پردازشی بر روی محتوای مالتی مدیا به‌کار گرفته می‌شود. در حین عملیات ویرایش بر روی محتوای مالتی مدیا جهت جعل نمودن، به‌طور مثال، در تصویر، ضروریست عملیات‌هایی همچون چرخش، هموارسازی، و بازنمونه‌برداری انجام شود. تشخیص حضور بازنمونه‌برداری تعیین می‌کند که آیا در تصویر دستکاری صورت گرفته است یا خیر [۲۵]. طرح کلی این روش در شکل (۲-۹) آورده شده است.

¹ Correlation based detection

² Spatial

³ Resampling

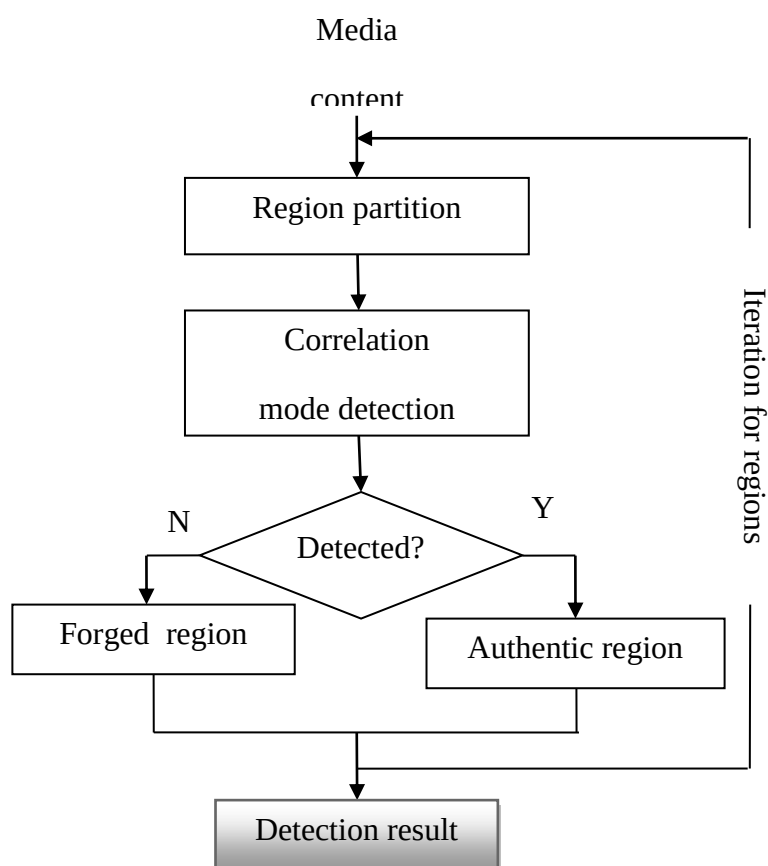
⁴ Color Filter Array (CFA) interpolation



شکل (۲-۹): تشخیص باز نمونه برداری [۵]

۲-۳-۱-۲- درونیابی آرایه‌ی فیلتر رنگی

آرایه‌ی فیلتر رنگی مولفه‌ای از دوربین‌های دیجیتال است. در زمان تولید تصویر دیجیتال، تنها یک سوم نمونه‌های تصویر رنگی توسط دوربین گرفته می‌شوند، در حالی که دو سوم باقیمانده توسط درونیابی آرایه فیلتر رنگی ایجاد می‌شود که باعث ایجاد همبستگی‌هایی در بین پیکسل‌های همسایه می‌شود. در خلال عملیات جعل تصویر، این همبستگی‌ها شاید از بین رفته و خراب شوند. بنابراین، تصدیق صحت تصویر، توسط آشکارسازی حضور یا عدم حضور ویژگی‌های درونیابی در تصویر، قابل تعیین است [۵]. طرح کلی این روش در شکل (۲-۱۰) آورده شده است.



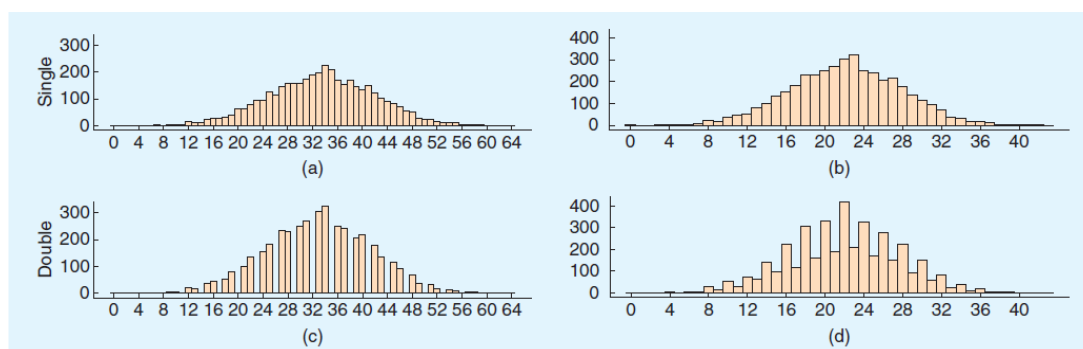
شکل (۲-۱۰): تشخیص جعل براساس روش درونیایی آرایه فیلتر رنگی [۵]

۲-۳-۳-۲- فشرده‌سازی JPEG مضاعف

در کدینگ JPEG، تصویر به بلوک‌هایی تقسیم‌بندی می‌شود. هر بلوک به وسیله‌ی تبدیل DCT^1 تبدیل می‌یابد، ضرائب DCT کوانتیزه شده و سپس با کدینگ آنتروپی کد می‌شوند. اگر تصویر، فشرده‌سازی JPEG مضاعف شود ضرائب DCT دوبار کوانتیزه می‌شود (معمولا نرم افزار ویرایش مدیا پس از جعل مدیا، آن را در یکی از قالب‌های فشرده سازی ذخیره می‌کند). کیفیت تصویر به وسیله گام کوانتیزاسیون کنترل می‌شود. تفاوت‌ها در دو گام کوانتیزاسیون عملیات مربوط را تغییر می‌دهند، که این امر، موجب اعوجاج‌هایی در تصویر نهایی می‌شود [۲۷]. اعوجاج ایجاد شده به وسیله کوانتیزاسیون

¹ Discrete Cosine Transform

مضاعف با اعوجاج ایجاد شده توسط یک کوانتیزاسیون متفاوت است، که این اختلاف را به وسیله‌ی روش ارائه شده در [۱۷] می‌توان مشخص نمود. با توجه به این روش، هیستوگرامی از ضرائب DCT بلوک‌های تصویر محاسبه می‌شود که سپس به وسیله تبدیل فوریه انتقال داده می‌شوند، که در نهایت منجر به تشخیص تناوب موقعیت‌های پیک در هیستوگرام تبدیل‌شده تصویر می‌شود [۲۸]. در حالت کلی، اگر تناوب تشخیص داده شود، کوانتیزاسیون مضاعف متناظر قابل تشخیص خواهد بود (شکل ۲-۱۱).



شکل (۲-۱۱): تشخیص جعل با توجه به فشرده‌سازی JPEG مضاعف. ردیف بالا هیستوگرام‌های سیگنال‌های کوانتیزه شده تک. شکل‌های نشان داده شده در ردیف پایین هیستوگرام‌های سیگنال‌های کوانتیزه شده مضاعف. به حاصل پررودیک مصنوعی در هیستوگرام‌های سیگنال‌های کوانتیزه شده مضاعف دقت نمایید [۲۷].

مولفان مقاله دقت تشخیص را بر روی بیش از ۱۰۰ تصویر مورد آزمایش قرار دادند. دقت قابل قبولی در تمام حالات بجز دو حالت گزارش شده است. در حالت اول، نسبت گام کوانتیزاسیون مرحله-ی دوم به گام مرحله‌ی اول یک مقدار صحیح است. این مسئله اعوجاج ایجاد شده به وسیله‌ی کوانتیزاسیون دوم را کاهش می‌دهد. در حالت دوم، گام اولین کوانتیزاسیون کوچک است، در حالی که گام مرحله‌ی دوم قابل توجه است. بنابراین، اولین کوانتیزاسیون در مقایسه با کوانتیزاسیون دوم قابل ملاحظه نیست. در هر دو حالت، کاهشی در میزان تناوب ایجاد شده در هیستوگرام ضرائب DCT مرتبط با فشرده‌سازی مضاعف وجود دارد.

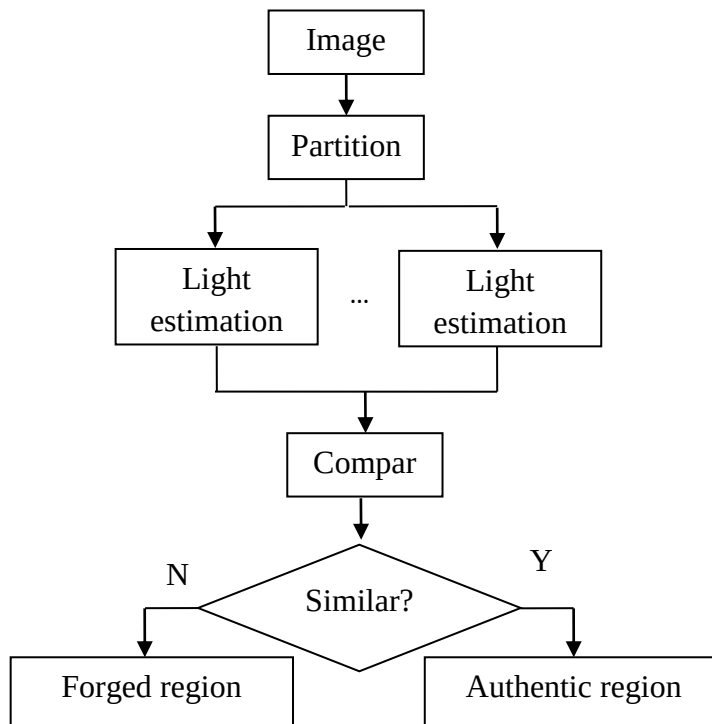
این روش آشکار می‌سازد که آیا یک تصویر دوبار فشرده‌سازی شده است یا خیر، که در صورت تایید این امر باعث به فعالیت پرداختن برای آشکار سازی نمونه‌های مختلف جعل می‌شود. بنابراین، در

یک تصویر با دوبار فشرده‌سازی، احتمال جعل افزایش می‌یابد. عیب این روش آسیب‌پذیری این روش در مقابل حملات است.

۲-۳-۳-۳- تشخیص برپایه‌ی ناهماهنگی در نور

تصور کنید که تنها یک نقطه نورانی در ارتباط با صحنه تصویر وجود دارد، پس شعاع‌های نورانی تخمین زده شده برای تمام اشیاء در تصویر باید با آن نقطه تقاطع پیدا کنند. به‌طور معکوس، دستکاری یک شیء در تصویر نشان خواهد داد که شعاع نور تخمین زده شده‌ی آن شیء با شعاع‌های نورانی اشیاء دیگر در همان تصویر در تناقض است. بنابراین، تصویر جعلی به‌وسیله تعیین مسیرهای نورانی قابل تشخیص است.

روش کلی برای استفاده از ناهماهنگی نوری در تشخیص جعل در [۲۹] شرح داده شده است. مولفان در ابتدا برخی روش‌های تشخیص جهت منبع نور را که در تحقیقات بینایی ماشین گزارش شده است معرفی کردند. این گزارش شامل تخمین منبع نورانی بی‌نهایت، منبع نوری محلی، یا چندین منابع نور است. سپس براساس تخمین، آن‌ها روش تشخیص را شرح دادند، همان‌گونه که در شکل (۲-۱۲) نشان داده شده است.



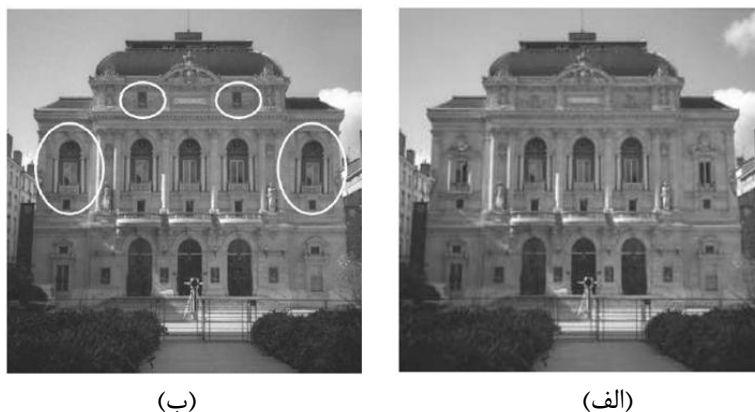
شکل (۲-۱۲): تشخیص جعل براساس ناهماهنگی نوری [۵]

جهات منبع نورانی در ارتباط با اشیای مختلف در تصویر با روش‌های مربوط تخمین زده می‌شوند؛ مقایسه‌ی مسیرهای نوری مختلف جعل یا اصل بودن تصویر را نشان می‌دهد.

۲-۳-۳-۴- تشخیص برپایه‌ی تیزی و مات شدگی تصویر

یکی از روش‌های ساده برای تشخیص عملیات‌های جعل و دستکاری در تصویر که شامل تغییر در تیزی/ مات شدگی تصویر است، تشخیص برپایه تیزی و مات شدگی تصویر است [۳۰]. این روش براساس این فرض است که اگر یک تصویر مورد جعل از نوع کپی- انتقال قرار بگیرد، میانگین مقدار تیزی/ مات شدگی در ناحیه جعل شده در قیاس با بخش‌های دست نخورده‌ی تصویر متفاوت است. روش تخمین میزان تیزی/ مات شدگی یک تصویر براساس خواص منظم ضرائب تبدیل ویولت که شامل اندازه‌گیری میزان اضمحلال ضرائب تبدیل ویولت در امتداد مقیاس‌ها است، می‌باشد. نتایج اولیه

نشان می‌دهد که مقدارهای کمی^۱ تیزی/ مات شدگی تخمین زده شده برای تشخیص نواحی دستکاری شده تصویر قابل استفاده است. تصویر در ابتدا به ناحیه‌هایی تقسیم‌بندی می‌شود، که در ادامه مقدار تیزی/ مات شدگی هر ناحیه محاسبه می‌شود. حضور مقدار تیزی/ مات شدگی در بعضی از قسمت‌ها که از یک مقدار آستانه توزیع نرمال تجاوز کرده باشد، برای تعیین ناحیه جعل شده استفاده می‌گردد. نتایج اولیه نشان می‌دهد که هنگامی که روش پیشنهادی به نواحی مختلف یک تصویر مات شده با میزان مشخص، اعمال می‌شود، به‌طور موفقیت آمیزی می‌تواند درجه‌ی مات شدگی تصویر مورد نظر را تخمین بزند. یکی از کاربردهای مفید برای چنین معیاری در حوزه مستندات دیجیتال قانونی است. تصویری که مورد جعل قرار گرفته است به احتمال بسیار زیاد اختلافاتی در ویژگی‌های مات شدگی یا تیزی آن وجود دارد. به منظور آزمایش طرح تشخیص جعل برپایه نظم و ترتیب ضرائب تبدیل ویولت، یک تصویر جعل شده دیجیتالی را مدنظر قرار داده و روش پیشنهادی به بخش‌های مختلف زیادی از آن تصویر اعمال شده است. تصویر اصلی و جعلی در شکل (۲-۱۳) نشان داده شده است.



شکل (۲-۱۳): نمونه‌ای از تشخیص جعل برپایه‌ی تیزی و مات شدگی تصویر. (الف)، تصویر اصلی و (ب)، تصویر جعل شده. بخش‌های جعل شده به‌وسیله دایره‌های روشن علامتگذاری شده‌اند. [۳۰]

۳۰ ناحیه‌ی متفاوت از نسخه جعلی تصویر مفروض بریده شده که ۴ تا از آنها جزء نواحی جعلی

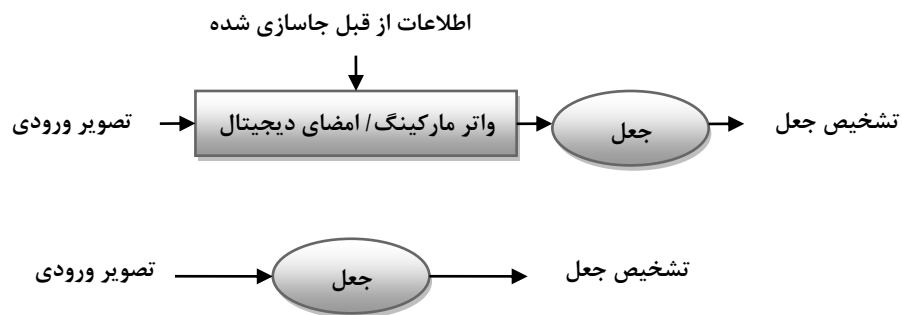
هستند.

^۱ Scores

فصل سوم

جعل کی - انتقال

روش‌های تشخیص جعل تصاویر دیجیتالی به دو دسته‌ی فعال^۱ و غیر فعال^۲ (یا کور^۳) تقسیم می‌شوند [۳۱] و [۳۲]. برخلاف روش‌های فعال مانند واتر مارکینگ دیجیتال و امضای دیجیتال، روش‌های غیر فعال نیازی به جاسازی اطلاعات در هنگام عکس برداری ندارند (شکل (۱-۳) را ببینید).



شکل (۱-۳): طرح کلی روش‌های فعال و غیر فعال. (بالا) تشخیص جعل فعال. (پایین) تشخیص جعل غیر فعال.

یکی از متداول‌ترین روش‌های جعل تصاویر، جعل کپی- انتقال است [۳۳] که در این تحقیق به بررسی روش‌های تشخیص آن خواهیم پرداخت. تکرار تصویر به معنی کپی کردن یک بخش از تصویر و چسباندن آن در یک موقعیت دیگر در همان تصویر و در نتیجه ایجاد دو بخش یکسان در تصویر می‌باشد. این نوع جعل به منظور افزایش تعداد اشیاء و یا پنهان کردن یک شیء یا ناحیه انجام می‌گیرد (شکل (۲-۳)) و به طور کلی به دو قسمت تقسیم می‌شود: روش‌های بر پایه‌ی بلوک بندی و روش‌های بر پایه‌ی نقاط کلیدی. در این فصل به بررسی این روش‌ها خواهیم پرداخت.

^۱ Active

^۲ Passive

^۳ Blind

در شکل زیر نمونه‌ای از این نوع جعل را برای پوشاندن شیء در تصویر مشاهده می‌کنید.



(ب)



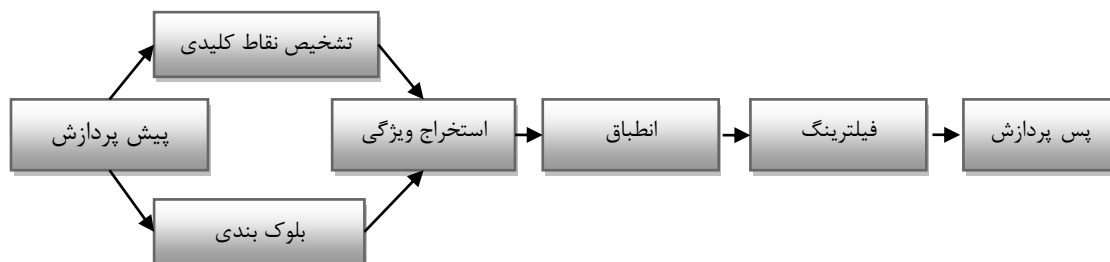
(الف)

شکل (۳-۲): نمونه‌ای از جعل کپی انتقال برای پوشاندن یک شیء. الف)، تصویر اصلی، ب) تصویر جعل شده برای پوشاندن تابلوی سمت راست.

۳-۱- فرآیند تشخیص جعل کپی- انتقال

روش‌های زیادی برای تشخیص جعل کپی-انتقال وجود دارد و بسیاری از آن‌ها برطبق روند

شکل (۳-۳) پیش می‌روند [۳۴].



شکل (۳-۳): فرآیند مشترک برای تشخیص جعل کپی-انتقال. استخراج ویژگی برای دو روش بر پایه‌ی نقاط کلیدی و بر پایه‌ی بلوک بندی متفاوت است. به غیر از مقادیر حد آستانه که برای هر روش متفاوت است، سایر مراحل مانند هم هستند.

روش‌های تشخیص جعل کپی- انتقال به طور کلی به دو دسته تقسیم می‌شوند: روش‌های بر

پایه‌ی بلوک بندی و روش‌های بر پایه‌ی نقاط کلیدی. که در ادامه به طور مفصل به بررسی این روش‌ها خواهیم پرداخت.

در هر دو روش بر پایه‌ی بلوک بندی ([۳]، [۵۹]–[۳۵]) بر پایه‌ی نقاط کلیدی ([۸]، [۶۰] و [۶۱]) بهتر است که تصویر پیش پردازش شود. به‌طور مثال، بیشتر روش‌ها بر روی تصاویر خاکستری عمل می‌کنند و برای این کار در ابتدا باید کانال‌های رنگی با هم ادغام شوند. برای استخراج ویژگی، روش‌های بر پایه‌ی بلوک بندی تصویر را به نواحی مستطیلی شکل تقسیم می‌کنند. برای هر کدام از آن نواحی، یک بردار ویژگی محاسبه می‌شود، سپس بردارهای ویژگی مشابه با هم انطباق می‌یابند. در مقابل، روش‌های بر پایه‌ی نقاط کلیدی، ویژگی‌ها را فقط در نواحی با آنتروپی بالا به‌دست می‌آورند، سپس ویژگی‌های مشابه تطبیق می‌یابند. در هر دو روش باید فیلترینگی برای از بین بردن انطباقات نادرست انجام شود. پس پردازش اختیاری نیز برای گروه بندی نقاط تطبیق یافته که مشترکاً از یک الگوی تبدیل پیروی می‌کنند، بر روی نواحی کشف شده انجام می‌شود. این روش‌ها عبارتند از:

الف) انطباق: در صورتی یک ناحیه را کپی شده تشخیص می‌دهیم که بین دو بردار توصیف کننده‌ی ویژگی‌شان، شباهت زیادی وجود داشته باشد. برای روش‌های بر پایه‌ی بلوک بندی، بسیاری از نویسنده‌ها، برای تعیین بردارهای ویژگی مشابه، استفاده از مرتب سازی بر مبنای واژه نگاری را پیشنهاد کرده‌اند. در مرتب سازی بر مبنای واژه نگاری، ماتریسی از بردارهای ویژگی طوری ساخته می‌شود که هر بردار ویژگی، در یک سطر از ماتریس قرار می‌گیرد. سپس ماتریس به صورت سطری مرتب می‌شود. بنابراین، ویژگی‌هایی با بیشترین شباهت در ردیف‌های متوالی ظاهر می‌شوند. روش مرتب سازی بر مبنای واژه نگاری، روش پیشنهادی این تحقیق، می‌باشد. سایر نویسنده‌ها از روش جست و جوی Best-Bin-First که از الگوریتم kd-tree مشتق شده است برای تقریب زدن نزدیک-ترین همسایه‌ها استفاده می‌کنند [۶۰]، [۶۲] و [۶۳]. در حقیقت روش‌های بر پایه‌ی نقاط کلیدی اغلب از این روش استفاده می‌کنند. انطباق با استفاده از روش kd-tree باعث جست و جوی موثرتری می‌شود. به طور نمونه از فاصله‌ی اقلیدسی به عنوان معیار شباهت استفاده می‌شود. در کارهای اولیه [۶۱]، نشان داده شده است که استفاده از روش انطباق kd-tree نتایج بهتری در مقابل مرتب سازی بر مبنای واژه نگاری دارد اما به حافظه‌ی بسیار بیشتری نیازمند است.

ب) **فیلترینگ:** پیکسل‌های همسایه، اغلب شدت^۱ یکسانی دارند که باعث اشتباه در تشخیص جعل می‌شوند. همچنین استفاده از معیارهای فاصله‌ی مختلف برای فیلتر کردن انطباقات ضعیف پیشنهاد شده است. به عنوان مثال، چندین نویسنده فاصله‌ی اقلیدسی را برای محاسبه‌ی فاصله‌ی بین بردارهای ویژگی پیشنهاد کردند [۶۳]، [۵۹]، [۳]. در مقابل، نویسنده‌ها در [7] استفاده از ضرایب همبستگی بین دو بردار ویژگی را به عنوان معیار شباهت پیشنهاد کردند.

ج) **پس پردازش:** هدف این مرحله‌ی آخر نگهداری انطباقاتی است که رفتار مشترکی را نشان می‌دهند. مجموعه‌ای از انطباقات را در نظر بگیرید که متعلق به یک ناحیه‌ی کپی شده است. انتظار می‌رود که این انطباقات از لحاظ مکانی در بلوک‌های مبدا و مقصد (یا نقاط کلیدی) بسیار به هم شبیه باشند. علاوه بر این انطباقاتی که حاصل عملیات کپی-انتقال یکسان هستند، دارای تعداد یکسانی از عملیات انتقال، تغییر سایز و چرخش هستند. نوعی از پس پردازش که بیشتر از همه استفاده می‌شود، نقاطی را که به اشتباه دسته بندی شده‌اند با تحمیل حداقل تعداد بردارهای انتقال بین نقاط تطبیق یافته تشخیص می‌دهد. یک بردار انتقال^۲ شامل انتقال (در مختصات تصویر) بین دو بردار ویژگی تطبیق یافته است. به عنوان مثال، دو بلوک از تصویر که کپی ساده‌ای (بدون چرخش یا تغییر سایز) از یکدیگر هستند را در نظر بگیرید، در این صورت هیستوگرام بردارهای انتقال، در پارامترهای انتقال عملیات کپی یک پیک^۳ را نشان می‌دهند.

۳-۲- روش‌های بر پایه‌ی بلوک بندی و روش‌های بر پایه‌ی نقاط کلیدی

روش‌های تشخیص جعل کپی-انتقال به طور کلی به دو دسته تقسیم می‌شوند: روش‌های بر پایه‌ی بلوک بندی و روش‌های بر پایه‌ی نقاط کلیدی. آزمایش‌ها نشان می‌دهد که روش‌های بر پایه‌ی

^۱ intensity

^۲ shift vector

^۳ peak

نقاط کلیدی مانند SIFT و SURF^۱، به خوبی ویژگی‌های بر پایه‌ی بلوک بندی مانند DCT، DWT، PCA، KPCA و زرنیک، عمل می‌کنند. این مجموعه ویژگی‌ها در برابر انواع مختلف نویزها مقاوم اند.

۳-۲-۱- الگوریتم‌های بر مبنای بلوک بندی

در این قسمت به بررسی ۱۳ ویژگی بر مبنای بلوک بندی می‌پردازیم که به چهار دسته تقسیم می‌شوند.

بر مبنای ممان^۲، بر مبنای کاهش بعد^۳، بر مبنای شدت^۴ و بر مبنای دامنه‌ی فرکانس (جدول (۱-۳) را ببینید).

جدول (۱-۳): گروه بندی مجموعه ویژگی‌ها برای تشخیص جعل کپی - انتقال

گروه	روش‌ها	طول ویژگی ^۵
ممان	BLUR [63]	۲۴
	HU [55]	۵
	ZERNIKE [53]	۱۲
کاهش بعد	PCA [52]	—
	SVD [43] KPCA [35]	— ۱۹۲
شدت	LUO [49]	۷
	BRAVO [38]	۴
	LIN [48]	۹
	CIRCLE [56]	۸
فرکانس	DCT [40]	۲۵۶
	DWT [35] FMT [37]	۲۵۶
		۴۵
نقاط کلیدی	SIFT [60], [61], [3]	۱۲۸
	SURF [62], [63]	۶۴

^۱ Speeded Up Robust Features

^۲ moment-based

^۳ dimensionality reduction- based

^۴ intensity-based

^۵ اندازه‌ی برخی از ویژگی‌ها بستگی به اندازه‌ی بلوک دارد، که معمولاً آن را ۱۶*۱۶ در نظر می‌گیریم. همچنین دقت کنید که اندازه‌ی ویژگی‌های PCA و SVD، به ترتیب بستگی به تصویر و محتوای بلوک دارد.

بر مبنای ممان: در این زمینه، سه روش مجزا را بررسی کرده‌ایم. در [۶۳] نویسنده استفاده از ممان‌های پایدار در برابر مات شدگی^۱ (BLURE) را به عنوان ویژگی پیشنهاد داده است. در [۵۵] نویسنده از چهار Hu مومن^۲ت اول به عنوان ویژگی استفاده کرده است (HU). در نهایت نویسنده ی [۵۳] استفاده از ممان‌های زرنیک (ZERNIKE) را پیشنهاد داده است.

بر مبنای کاهش بعد: در [۵۲]، فضای برداهای تطبیق یافته از طریق PCA کاهش پیدا می-کند. نویسنده ی [۳۵]، استفاده از KPCA^۲ (نوعی از PCA) را پیشنهاد کرده است. در [۴۳] برای کاهش بعد استفاده از SVD^۳ پیشنهاد شده است. چهارمین روش، ترکیبی از محاسبه ی DWT و SVD می‌باشد [۴۵]. البته این روش نتایج قابل اعتمادی به‌دست نمی‌دهد و از ارزیابی حذف شده است.

بر مبنای شدت: اولین سه ویژگی که در [۴۹] و [۳۸] استفاده شده، میانگین اجزای قرمز، سبز و آبی است. علاوه بر این در [۴۹] استفاده از اطلاعات هدایت کننده ی بلوک‌ها (LUO) پیشنهاد شده است. در حالی که نویسنده در [۳۸] آنتروپی بلوک‌ها را به عنوان ویژگی‌های تمایز دهنده در نظر گرفته است (BRAVO). در [۴۸] میانگین شدت‌های خاکستری یک بلوک و زیر بلوک‌های محاسبه می‌شود. در [۵۶] از میانگین شدت‌های دایره‌ها با شعاع‌های مختلف در اطراف مرکز بلوک استفاده شده است (CIRCLE).

بر مبنای فرکانس: در [۴۰] استفاده از ۲۵۶ ضرایب تبدیل کسینوسی گسسته به عنوان ویژگی پیشنهاد شده است (DCT). ضرایب تبدیل ویولت گسسته^۴ (DWT) با استفاده از ویولت‌های Haar، به عنوان ویژگی در [۳۵] پیشنهاد شده است. نویسنده در [۳۷] استفاده از FMT^۵ را برای ایجاد بردار ویژگی پیشنهاد کرده است. در فصل بعدی به طور مفصل به بررسی این روش‌ها و

^۱ blur-invariant moments

^۲ Kernel-PCA

^۳ singular value decomposition

^۴ Discrete Wavelet Transform

^۵ Fourier-Mellin Transform

ویژگی‌ها پرداخته و مقایسه‌ای بین تمام روش‌های ذکر شده خواهیم داشت.

۳-۲-۲- الگوریتم‌های بر مبنای نقاط کلیدی

بر خلاف الگوریتم‌های بر مبنای بلوک بندی، روش‌های بر مبنای نقاط کلیدی، بر پایه‌ی شناسایی و انتخاب نواحی با آنترپی بالا در تصویر است (یعنی "نقاط کلیدی"). سپس به ازای هر نقطه‌ی کلیدی، یک بردار ویژگی استخراج می‌شود. در نتیجه تعداد کمتری بردار ویژگی تخمین زده می‌شود که باعث کاهش پیچیدگی محاسباتی در انطباق ویژگی‌ها و پس پردازش می‌شود. تعداد کمتر بردارهای ویژگی باعث می‌شود که مقدار حد آستانه‌های عملیات پس پردازش کمتر از روش‌های بر مبنای بلوک بندی باشد. یکی از مشکلات روش‌های بر مبنای نقاط کلیدی این است که گاهی اوقات در نواحی کپی شده، نقاط کلیدی به صورت پراکنده و کم قرار دارند. اگر مناطق کپی شده تقریباً هموار باشند، توسط این روش قابل تشخیص نیستند. دو نوع مختلف از بردارهای ویژگی بر مبنای نقاط کلیدی را بررسی کردیم. یکی از محققین از ویژگی‌های SIFT و دیگری از ویژگی‌های SURF استفاده کرده‌اند [۶۲]. این روش‌ها با SIFT و SURF نشان داده شده‌اند. تفاوت روش‌های استخراج ویژگی بر مبنای نقاط کلیدی در پس پردازش ویژگی‌های تطبیق یافته است.

فصل چهارم

بررسی کارهای انجام شده در زمینه‌ی تشخیص جعل

کپی – انتقال

در سال‌های اخیر، محققان تلاش‌های زیادی در راستای تشخیص جعل داده‌های مالتی مدیا انجام داده‌اند. همان‌طور که گفتیم، یکی از معمول‌ترین روش‌های جعل، جعل کپی-انتقال است. هدف در تشخیص جعل کپی-انتقال، تعیین نواحی کپی شده است حتی اگر کمی با هم متفاوت باشند. اکثر جاعلان برای این‌که تصویر طبیعی‌تر به نظر برسد، تغییراتی مانند افزودن نویز، مات کردن، دستکاری کردن^۱، فشرده سازی و... همچنین تبدیلات هندسی مانند تغییر سایز و چرخش را بر روی ناحیه‌ی کپی شده اعمال می‌کنند. برخی ویژگی‌ها که از تصویر استخراج می‌شوند، در برابر تغییرات و تبدیلات هندسی مقاوم نیستند، بنابراین، استفاده از این ویژگی‌ها در تشخیص جعل زمانی مناسب است که هیچ تغییری بر روی ناحیه‌ی کپی شده اعمال نشده باشد. برخی دیگر از آن‌ها در برابر انواع تبدیلات و تغییرات مقاوم هستند و استفاده از آن‌ها در این مواقع، مقاومت و دقت تشخیص را بالا می‌برد. در این فصل به بررسی روش‌های تشخیص جعل که تا این زمان توسط محققان پیشنهاد شده و مزایا و معایب هر کدام از آن‌ها خواهیم پرداخت.

۴-۱- بررسی مقالات و کارهای انجام شده

[۶۴] به بررسی انواع روش‌های موجود در تشخیص جعل کپی انتقال پرداخته است. راه حل مستقیم برای تشخیص این نوع جعل، استفاده از جستجوی جامع^۲ است. این روش ممکن است که از لحاظ محاسباتی بسیار پرهزینه باشد، و برای یک تصویر $M \times N$ هزینه‌ای حدود $(MN)^2$ دارد. از طرفی این روش در مواردی که ناحیه‌ی کپی شده دچار تغییراتی شده باشد، به درستی عمل نمی‌کند. روش دوم که موثرتر از روش اول است و در [۱] آمده است، استفاده از ویژگی‌های خود همبستگی^۳ است. با این حال، این روش زمانی مؤثر است که بخش کپی شده، بخش بزرگی از تصویر باشد.

^۱ retouching

^۲ exhaustive search

^۳ auto correlation

روش دیگر، فرآیند تطبیق بلوک است. در این روش، ابتدا تصویر به بلوک‌های هم پوشانی تقسیم می‌شود. در این روش می‌خواهیم به جای تعیین تمام ناحیه‌ی کپی شده، بلوک‌های متصل به هم از تصویر را که کپی شده‌ی یکدیگر هستند پیدا کنیم. توجه کنید که ناحیه‌ی کپی شده شامل تعداد زیادی از بلوک‌های هم پوشانی است و از آنجایی که هر بلوک با شیفت یکسانی انتقال داده شده است، فاصله‌ی بین هر زوج بلوک کپی شده یکسان است. بنابر این زمانی جعل تشخیص داده می‌شود که بیشتر از یک تعداد مشخص از بلوک‌های تصویر با فواصل یکسان وجود داشته باشند و این بلوک‌ها طوری به هم متصل باشند که تشکیل دو ناحیه‌ی یک شکل را بدهند.

نکته‌ی مهم در این روش‌ها این است که راهی مناسب برای نمایش بلوک‌ها پیدا کنیم به طوریکه بلوک‌های کپی شده، حتی اگر دستخوش تغییرات قرار گرفته باشند هم قابل تشخیص باشند. برخی از نویسندگان استفاده از ویژگی‌های مختلف را برای نمایش بلوک‌ها پیشنهاد کرده اند. این ویژگی‌ها و برخی خصوصیات آن‌ها در بخش بعدی بحث خواهد شد. مسئله‌ی دیگر این است که بتوانیم جفت بلوک‌هایی را که نمایش یکسانی دارند، پیدا کنیم.

۴-۲- ویژگی‌ها و مقایسه‌ی آن‌ها در جعل کپی - انتقال

برای تشخیص جعل بر اساس بلوک بندی، ابتدا تصویر را به بلوک‌های هم پوشانی تقسیم می‌کنیم. سپس عملیات استخراج ویژگی را انجام می‌دهیم. ویژگی‌ها باید طوری تعیین شوند که برای بلوک‌های کپی شده، حتی اگر دچار تغییرات شده باشند، مقادیر یکسانی داشته باشند.

نویسندگان در [۱] و [۶۵] برای این منظور از ضرایب DCT استفاده کرده‌اند. مزیت DCT این است که انرژی سیگنال بر روی چند ضریب اول متمرکز است و سایر ضرایب بسیار کوچک هستند. بنابراین، تغییراتی که در تصویر در فرکانس‌های بالا رخ می‌دهد، مانند افزودن نویز، فشرده سازی و دست‌کاری کردن تصویر، بر روی این ضرایب ابتدایی تأثیری نمی‌گذارد. نویسندگان در این مقاله مقاوم

بودن روش خود را در برابر انواع دستکاری کردن تصویر آزمایش کردند، ولی آزمایشی مبنی بر مقاوم بودن روششان در برابر سایر تغییرات انجام نداده‌اند.

سپس محققین استفاده از PCA را به عنوان جایگزینی برای DCT برای نمایش بلوک‌ها پیشنهاد کردند و نتایج کار خود را در [۵۲] منتشر کردند. برای این منظور، ضرایب در هر بلوک، محاسبه کردند و به بردار تبدیل کردند، آنگاه این بردارها را در یک ماتریس به صورت سطری قرار دادند و ماتریس کواریانس^۱ متناظر آن را محاسبه کردند. برای کاهش بعد، تصویر هر بلوک بر روی این بردارهای پایه با مقادیر ویژه^۲ بزرگ‌تر به دست آمد و به عنوان نمایشی جدید از بلوک‌ها از آن‌ها استفاده شد. این روش نمایش، نسبت به فشرده سازی و افزودن نویز مقاوم است [۶۶]، ولی دوباره نمونه^۳ برداری از تصویر مانند تغییر سایز و چرخش، بر روی این مقادیر ویژه تأثیر خواهد گذاشت. نشان داده شده است که روش آن‌ها در برابر فشرده سازی JPEG تا سطح ۵۰ و نویز با SNR ای برابر با ۳۶dB و ۲۹dB مقاوم است.

در مقالات جدیدتر، به جای استخراج ویژگی، به صورت مستقیم از تصویر، پیشنهاد شده است که ابتدا تصویر را با استفاده از DWT به چهار زیر باند تقسیم کنیم [۴۵]. برای کاهش تعداد بلوک‌ها و افزایش سرعت محاسبات، با توجه به این که بیشترین انرژی در این زیر باند متمرکز شده است، زیر باند با فرکانس پایین را به بلوک‌های هم پوشانی تقسیم کردند. سپس SVD را بر روی این بلوک‌ها اعمال کردند. آن‌ها نشان دادند که روششان نسبت به فشرده سازی JPEG تا سطح کیفی ۷۰ مقاوم است.

در [۳۹] نویسندگان تکنیکی ارائه داده که در آن اگر ناحیه چسبانده شده توسط دو ابزار Adobe Photoshop healing brush و Poisson cloning در فتوشاپ دستکاری شود، قابل تشخیص است.

^۱ covariance matrix

^۲ eigenvalue

^۳ re-sampling

در [۴۹] نویسندگان استفاده از ویژگی‌هایی بر اساس اطلاعات رنگی بلوک‌ها را پیشنهاد کرده است. اولین مجموعه از ویژگی‌ها شامل میانگین اجزای رنگی رنگ‌های قرمز، سبز و آبی است. در دومین مجموعه، بلوک‌ها به دو قسمت و در چهار جهت تقسیم می‌شوند. نسبت مقدار شدت یک قسمت به شدت کل بلوک محاسبه می‌شود. نتایج نشان می‌دهد که این روش نسبت به فشرده سازی JPEG تا سطح کیفی ۳۰ و همچنین در برابر گوسین بلورینگ^۱ (پنجره ۵ × ۵ و $\sigma=1$) و نویز افزایشی نیز مقاوم می‌باشد.

در [۳۷] و [۶۷] نویسندگان FMT را بر روی بلوک‌های تصویر اعمال کرد. برای این منظور ابتدا نمایش تبدیل فوریه‌ی هر بلوک را به دست آوردند و سپس مقادیر حاصل را به مختصات قطبی بردند. سپس نمایش برداری را با تصویر کردن مقادیر مختصات قطبی بر روی یک بعد به دست آوردیم و از آن‌ها به عنوان ویژگی استفاده کردیم.

یک روش جعل خوب باید نسبت به تبدیلاتی مانند تغییر سائز و چرخش و همچنین نسبت به دستکاری‌هایی مانند فشرده سازی JPEG، نویز افزایشی گوسین، و اصلاح گاما مقاوم باشد. اغلب روش‌های موجود نسبت به تمام این موارد مقاوم نیستند و اغلب از لحاظ محاسباتی پیچیده و زمانبر هستند. به عنوان مثال، روش ارائه شده در [۵۲] نسبت به چرخش و تغییر سائز مقاوم نیست و روش‌های [۱] و [۳۷] تنها تغییرات کوچکی در چرخش و تغییر سائز را تشخیص می‌دهند.

نویسندگان در [۳]، با استفاده از ممان‌های زرنیک به تشخیص جعل زمانیکه تنها عملیات چرخش بر روی ناحیه کپی شده اعمال شده است می‌پردازند. این روش، در تشخیص جعل در نواحی هموار نیز موفق بوده است.

^۱ Gaussian blurring

امروزه از ویژگی‌های بصری محلی^۱ (مانند SIFT، SURF، GLOH و...) به دلیل مقاومتشان در برابر تبدیلات هندسی (مانند چرخش و تغییر سایز) به طور وسیعی در بازیابی تصویر^۲ و تشخیص شیء استفاده می‌شود.

در سال ۲۰۰۶م. یک توصیف‌گر ویژگی غیر حساس به تغییر سایز دیگر به نام SURF [۶۸] ارائه شد، این توصیف‌گر همانند توصیف‌گر SIFT نقاط کلیدی را استخراج می‌کند با این تفاوت که هر نقطه توسط یک بردار ویژگی با ۶۴ مؤلفه نمایش داده می‌شود. بر اساس ارزیابی‌های انجام شده، محققان به این نتیجه رسیده‌اند که SIFT قدرت تشخیص بیشتری نسبت به SURF دارد ولی، در مقابل، SURF سریع‌تر و منسجم‌تر عمل می‌کند.

اخیراً از ویژگی‌های SIFT در زمینه‌ی دیجیتال قانونی استفاده می‌شود. در حقیقت ویژگی‌های SIFT در تشخیص اثر انگشت [۶۹] و همچنین برای تشخیص جعل کپی-انتقال استفاده می-شود [۷۰]، [۶۰]، [۴].

در این پایان نامه به معرفی روشی ترکیبی برای تشخیص جعل کپی-انتقال، با استفاده از ویژگی‌های SIFT و ممان‌های زرنیک می‌پردازیم. در نهایت در [۶۱] و [۷۱] مقایسه‌ای بین روش‌های ذکر شده و ارزیابی عملکرد هر کدام از این روش‌ها انجام گرفته است.

^۱ local visual features

^۲ image retrieval

فصل پنجم

ویژگی‌های SIFT و ممان‌های زرنیک

همان‌طور که گفته شد، روش‌های تشخیص جعل کپی - انتقال به طور کلی به دو دسته تقسیم می‌شوند: روش‌های بر پایه‌ی بلوک بندی و روش‌های بر پایه‌ی نقاط کلیدی. آزمایش‌ها نشان می‌دهد که روش‌های بر پایه‌ی نقاط کلیدی مانند SIFT و SURF، به خوبی ویژگی‌های بر پایه‌ی بلوک بندی مانند DCT، DWT، KPCA، PCA و زرنیک، عمل می‌کنند. این مجموعه ویژگی‌ها در برابر انواع مختلف نویزها مقاوم اند.

در این فصل به بررسی خواص و روابط ریاضیاتی ویژگی‌های SIFT و ممان‌های زرنیک و نحوه‌ی استخراج آن‌ها از تصاویر می‌پردازیم.

۵-۱- ویژگی‌های مقاوم در برابر تغییر سائز (SIFT)

الگوریتم SIFT به وسیله‌ی David Lowe در سال ۱۹۹۴ ابداع گردید. SIFT یک روش برای آشکارسازی و استخراج ویژگی‌های مستقل و مشخص از تصاویری می‌باشد که می‌تواند برای تناظریابی بین تصاویر و یا شیء با یک تصویر مورد استفاده قرار گیرد [۷۲].

مراحل اصلی آشکار کردن و استخراج نقاط کلیدی SIFT در چهار مرحله خلاصه می‌شود:

۵-۱-۱- تعیین نقاط مینیمم و ماکزیمم در فضای مقیاس

گام نخست برای یافتن نقاط کلیدی، پیدا کردن مکان‌هایی است که در برابر تغییر مقیاس تصویر ثابت هستند. این نقاط با استفاده از جستجو برای ویژگی‌های ثابت در تمامی مقیاس‌های ممکن، توسط تابع مقیاس که بنام فضای مقیاس شناخته شده است به دست می‌آیند.

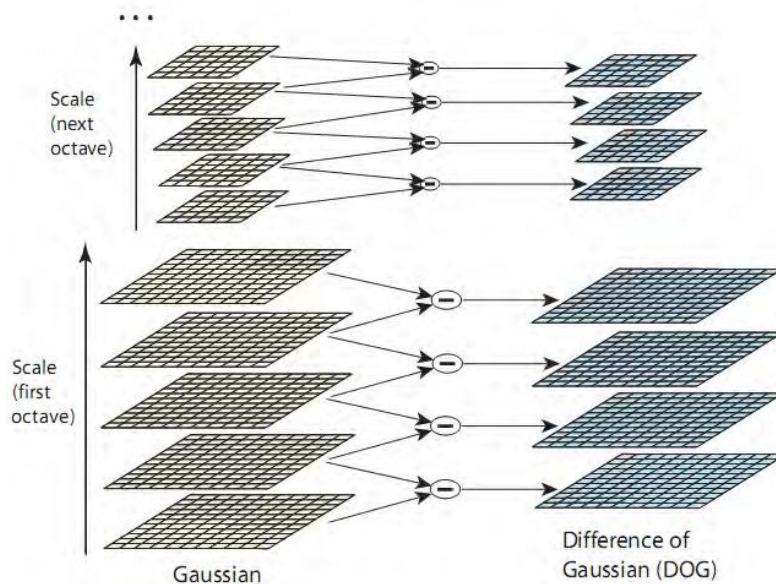
این کار با استفاده از مفاهیم هرم تصویری و فضای مقیاس انجام می‌پذیرد. به وسیله‌ی هرم تصویری، با توجه به مفهوم فضای مقیاس که در آن، تصویر در مقیاس‌های مختلف ساخته می‌شود معمولاً یک تصویر دیجیتالی در قدرت تفکیک‌های متفاوت نمایش داده می‌شود. می‌توان مطابق شکل (۱-۵) ابتدا تصویر را در چند مقیاس مختلف تشکیل داد و سپس برای هر مقیاس تصویری، پنج تصویر (اکتاو) با انحراف معیار مختلف به وجود آورد که در نتیجه به تعداد مقیاس‌های مختلف دسته‌های پنج تصویری خواهیم داشت. با توجه به این که از خصوصیات اصلی تصاویر هرمی کاهش نویز تصویر در اثر هموار سازی در قدرت تفکیک‌های پایین تر می‌باشد، لذا در این روش از فیلتر گوسین استفاده شده است و دسته‌های پنج‌تایی با انحراف معیارهای مختلف توسط فیلتر گوسین و ضرب کرنل گوسین در تصویر اصلی با انحراف معیارهای متفاوت به وجود آمده است. سپس همان گونه که در شکل (۱-۵) ملاحظه می‌گردد، هر دو تصویر مجاور در داخل دسته‌های پنج‌تایی را با اختلاف تفاضلی بر اساس جبر ماتریسی از هم کم کرده و تصاویر جدید را بر اساس اختلاف گوسین (DOG^1) ایجاد می‌نماییم. این مشتق گیری ثانویه در واقع برای پیدا کردن نقاط کمینه و بیشینه در فضای مقیاس صورت می‌پذیرد و در نتیجه مقادیر کمینه و بیشینه در تصاویر بطور خاص در مقیاس‌های مختلف ایجاد می‌گردد. در این مرحله، ویژگی‌ها مستقل از مقیاس می‌گردند. معادلات (۱-۵)، معادلات انجام این مرحله است که در آن I تصویر اصلی و G تابع کرنل گوسین می‌باشد که با ضریب تلفیقی (کانولوشن)؛ تصویر هموار شده L را تولید می‌نماید، سپس با کم کردن دو تصویر حاصل شده‌ی L ، در هر دسته پنج‌تایی با تصویر مجاورش که اختلاف آن‌ها در انحراف معیار σ می‌باشد، تصویر D بر اساس اختلاف تفاضلی گوسین ایجاد می‌گردد.

¹ difference of gaussian

$$L(x, y, \sigma) = G(x, y, \sigma) \times I(x, y)$$

$$D(x, y, \sigma) = (G(x, y, k\sigma) - G(x, y, \sigma)) \times I(x, y)$$

$$G(x, y, \sigma) = \frac{1}{2\pi\sigma^2} \exp \left(-\frac{(x^2 + y^2)}{2\sigma^2} \right) \quad (1-5)$$



شکل (۱-۵): هرم تصویری با تابع گوسین [۱]

۵-۱-۲- انتخاب نقاط کلیدی از بین نقاط ماکزیمم و مینیمم اصلی

در این مرحله مقادیر شدت خاکستری هر پیکسل با ۸ پیکسل مجاورش و با ۹ پیکسل در تصاویر مجاور بالایی و پایینی که از لحاظ σ با هم اختلاف داشتند، مقایسه می‌گردد. اگر مقدار این پیکسل از تمام ۲۶ پیکسل همسایه بیشتر یا کمتر بود آن را به عنوان یک نقطه‌ی کاندیدا انتخاب و آن را نقطه‌ی کلیدی می‌نامیم.

۵-۱-۳- تعیین موقعیت مکانی هر نقطه‌ی کلیدی

در این مرحله به‌طور دقیق موقعیت هر نقطه کلیدی از لحاظ مختصات تعیین می‌گردد و نقاطی که دارای کنتراست کم می‌باشند حذف می‌شوند، سپس مقدار شیب (ضریب زاویه‌ای) و همچنین زاویه‌ی شیب هر نقطه‌ی کلیدی تعیین می‌گردد.

$$m(x, y) = \sqrt{(L(x+1, y) - L(x-1, y))^2 + (L(x, y+1) - L(x, y-1))^2}$$

$$\theta(x, y) = \tan^{-1} \left(\frac{L(x, y+1) - L(x, y-1)}{L(x+1, y) - L(x-1, y)} \right) \quad (۲-۵)$$

در نتیجه برای تمام نقاط اصلی ماتریس $n \times 4$ ، به صورت زیر تشکیل می‌گردد.

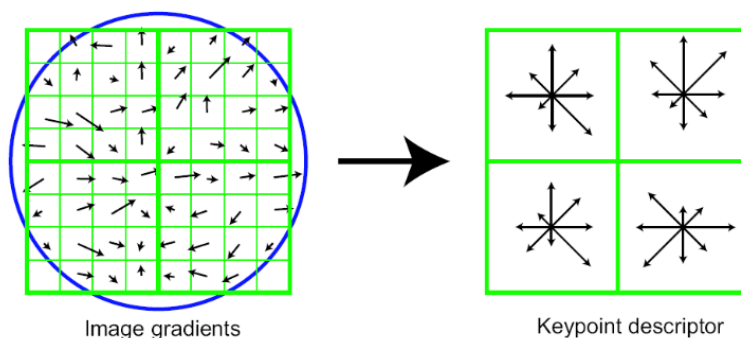
$$\begin{bmatrix} x_1 & y_1 & m_1 & \theta_1 \\ 0 & 0 & 0 & 0 \\ x_n & y_n & m_n & \theta_n \end{bmatrix}_{n \times 4} \quad (۳-۵)$$

۵-۱-۴- تشکیل توصیف کننده برای هر نقطه‌ی کلیدی

به ازای هر نقطه کلیدی، پیکسل‌های مجاور آن (4×4) انتخاب می‌شوند و در هر پیکسل ضریب زاویه‌ای (گرادیانت) تعیین و در مرکز ۱۶ پیکسل (4×4) برآیند هر یک بر روی چهار جهت اصلی و نیمساز آن ترسیم می‌شود و تشکیل هیستوگرام می‌دهد. هر هیستوگرام شامل ۸ شاخه و هر توصیف-گر شامل یک آرایه‌ی ۴ هیستوگرامی خواهد شد. در نتیجه یک بردار ویژگی SIFT شامل یک بردار با $4 \times 4 \times 8 = 128$ عنصر خواهد بود که با استفاده از نرمال سازی آن هر توصیف‌گر عدم وابستگی به تغییرات روشنایی نیز محقق خواهد شد.

بردار ویژگی‌ها برای کاهش اثر تغییرات روشنایی اصلاح می‌شوند. در ابتدا بردار به طول واحد نرمال می‌شود. در مورد تغییر در کنتراست تصویر، مقدار هر پیکسل در ثابتی ضرب می‌شود، گرادیان

نیز در همان ثابت ضرب می‌شود. لذا تغییر در کنتراست با نرمال سازی بردار برطرف خواهد شد. برای مقاوم سازی در برابر تغییرات روشنایی، که یک ثابت به هر پیکسل اضافه می‌شود، تغییری در مقدار گرادیان نخواهد داشت. بنابراین، توصیف‌گر در مقابل تغییرات افاین و روشنایی مقاوم است. برای کاهش تأثیر نامطلوب دامنه‌های بزرگ گرادیان، آن‌ها را با آستانه‌گذاری مناسب محدود به $0/2$ می‌کنیم و سپس به مقدار واحد تبدیل می‌نماییم. این بدین معناست که تطبیق دامنه‌ها برای نقاط با گرادیان‌های بزرگ اهمیتی ندارند در حالی که توزیع جهت‌ها اهمیت بیشتری دارد. مقدار $0/2$ نیز به صورت تجربی با استفاده از تصاویری با روشنایی متفاوت برای تصویر یک شیء سه بعدی به دست آمده است. در نهایت، یک بردار 128 عضوی در توصیف‌گر قرار می‌گیرد. در مرحله‌ی تشخیص اشیاء، ابتدا نقاط کلیدی تصاویر مدل و آزمون (تصاویری که می‌خواهیم بر هم منطبق کنیم) را استخراج کرده، اقدام به تطبیق زوج نقاط متناظر در دو تصویر می‌کنیم. به دلیل وجود نویز در ویژگی‌های استخراج شده، تطبیق‌های اولیه نقاط کلیدی ممکن است نادرست باشد.



شکل (۵-۲): توصیف کننده‌های محلی در تصویر برای یک نقطه‌ی کلیدی

۵-۲- ممان‌های زرنیک

در زمینه‌ی شناسایی الگو، کاربردهای واترمارک دیجیتال و ...؛ مومنت‌ها و توابع آن‌ها، برای استخراج ویژگی‌های تغییر ناپذیر به طور وسیع استفاده می‌شوند [۷۳] و [۷۴]. از بین انواع مختلف ممان‌ها، ممان‌های زرنیک [۳] به خاطر حساس نبودن به نویز و محتوای تصویر، مورد توجه بیشتری

قرار گرفته‌اند [۷۵] و [۷۶]. ممان‌های زرنیک، در برابر چرخش، ثابت هستند. روش‌های تشخیص جعل بر مبنای ممان‌های زرنیک، در برابر تغییرات نویز گوسین سفید افزایشی، فشرده سازی JPEG و مات شدگی مقاوم است. در این قسمت به بررسی ممان‌های زرنیک و بیان روابط ریاضیاتی آن‌ها پرداخته‌ایم و اصول کلی بر اساس [۷۵] و [۷۶] می‌باشد.

۵-۲-۱- تعریف

ممان‌های زرنیک مرتبه‌ی n با تکرار m ، برای یک تابع پیوسته‌ی تصویر، $f(x, y)$ که در خارج از یک دایره‌ی واحد به صفر می‌رسند، به صورت زیر هستند [۷۷]:

$$A_{nm} = \frac{n+1}{\pi} \iint_{x^2+y^2 \leq 1} f(x, y) V_{nm}^*(\rho, \theta) dx dy, \quad (۴-۵)$$

که n یک عدد صحیح غیر منفی و m یک عدد صحیح است به طوریکه $n-|m|$ غیر منفی و زوج است. تابع مقادیر مختلط، $V_{nm}(x, y)$ به صورت زیر تعریف می‌شوند:

$$V_{nm}(x, y) = V_{nm}(\rho, \theta) = R_{nm} \exp(jm\theta) \quad (۵-۵)$$

که ρ و θ مختصات قطبی بر روی یک دایره‌ی واحد نشان می‌دهند. و R_{nm} ، چند جمله‌ای ρ هستند (چند جمله‌ای‌های زرنیک) و به صورت زیر نشان داده می‌شوند:

$$R_{nm}(\rho) = \sum_{s=0}^{(n-|m|)/2} \frac{(-1)^s [(n-s)!] \rho^{n-2s}}{s! (\frac{n+|m|}{2} - s)! (\frac{n-|m|}{2} - s)!}. \quad (۶-۵)$$

توجه کنید که $R_{n,-m}(\rho) = R_{nm}(\rho)$. این چند جمله‌ای ها متعامد هستند :

$$\iint_{x^2+y^2 \leq 1} [V^*(x,y)] \times V_{pq}(x,y) dx dy = \frac{\pi}{n+1} \delta_{np} \delta_{mq}, \quad (7-5)$$

$$\delta_{ab} = \begin{cases} 1, & a=b \\ 0, & \text{otherwise} \end{cases}$$

برای تصاویر دیجیتال، در رابطه‌ی بالا، انتگرال‌ها با سیگما جایگزین می‌شوند. برای محاسبه‌ی ممان زرنیک یک بلوک داده شده، مرکز بلوک به عنوان مبدأ در نظر گرفته می‌شود و مختصات پیکسل‌ها در محدوده‌ی دایره‌ی واحد نگاشت می‌شوند. پیکسل‌هایی که در خارج از دایره قرار می‌گیرند در محاسبات استفاده نمی‌شوند. توجه کنید که $A_{nm}^* = A_{n,-m}$.

فرض کنید تمام مومنت‌های A_{nm} تابع $f(x, y)$ تا مرتبه‌ی $n_{\max}$ را داریم. تابع گسسته‌ی تصویر اصلی $\hat{f}(x, y)$ که مومنت‌هایش همان مومنت‌های $f(x, y)$ تا مرتبه‌ی $n_{\max}$ است، می‌تواند محاسبه شود. با استفاده از تعامد زرنیک‌های پایه می‌توانیم $\hat{f}(x, y)$ را به صورت زیر از نو بسازیم:

$$\hat{f}(x, y) = \sum_{n=0}^{n_{\max}} \sum_m A_{nm} V_{nm}(\rho, \theta). \quad (8-5)$$

توجه کنید که اگر $n_{\max}$ به سمت بی‌نهایت برود، $\hat{f}(x, y)$ به سمت $f(x, y)$ میل می‌کند.

۵-۲-۲- مقاوم بودن ممان‌های زرنیک در برابر چرخش

در این قسمت به اثبات مقاوم بودن ممان‌های زرنیک در برابر چرخش می‌پردازیم. فرض کنید که تصویر به اندازه‌ی α چرخیده است. اگر تصویر چرخیده را با f' نمایش دهیم، ارتباط بین تصویر اصلی و تصویر چرخیده در مختصات قطبی یکسان به صورت زیر است:

$$f'(\rho, \theta) = f(\rho, \theta - \alpha). \quad (9-5)$$

با استفاده از معادلات (4-5) و (5-5)، داریم:

$$\begin{aligned} A_{nm} &= \frac{n+1}{\pi} \int_0^{2\pi} \int_0^1 f(\rho, \theta) Y_{nm}^*(\rho, \theta) \rho d\rho d\theta \\ &= \frac{n+1}{\pi} \int_0^{2\pi} \int_0^1 f(\rho, \theta) R_{nm}(\rho) \exp(-jm\theta) \rho d\rho d\theta. \end{aligned} \quad (10-5)$$

بنابراین، ممان‌های زرنیک تصویر چرخیده در مختصات یکسان به صورت زیر تعریف می‌شوند:

$$A'_{nm} = \frac{n+1}{\pi} \int_0^{2\pi} \int_0^1 f(\rho, \theta - \alpha) R_{nm}(\rho) \exp(-jm\theta) \rho d\rho d\theta. \quad (11-5)$$

با یک تغییر متغیر به صورت $\theta_1 = \theta - \alpha$ خواهیم داشت:

$$\begin{aligned} A'_{nm} &= \frac{n+1}{\pi} \int_0^{2\pi} \int_0^1 f(\rho, \theta_1) R_{nm}(\rho) \exp(-jm(\theta_1 + \alpha)) \rho d\rho d\theta_1 \\ &= \left[\frac{n+1}{\pi} \int_0^{2\pi} \int_0^1 f(\rho, \theta) R_{nm}(\rho) \exp(-jm\theta) \rho d\rho d\theta_1 \right] \exp(-jm\alpha) \\ &= A_{nm} \exp(-jm\alpha). \end{aligned} \quad (12-5)$$

معادله‌ی (12-5) نشان می‌دهد که هر ممان زرنیک در اثر چرخش، دچار تغییر زاویه‌ی فاز می‌شود. بنابراین، مقدار ممان زرنیک، $|A_{nm}|$ می‌تواند به عنوان ویژگی ثابت در برابر چرخش در تصویر استفاده شود. بنابراین، مقدار ممان‌های زرنیک را برای توصیف منحصر به فرد هر بلوک، بدون در نظر گرفتن چرخش محاسبه می‌کنیم.

فصل ششم

روش پیشنهادی و مقایسه‌ی آن با روش‌های دیگر

در این تحقیق به طور خاص به بررسی و تشخیص جعل کپی - چسباندن یا کپی - انتقال می - پردازیم. اکثر تصاویر جعلی با چشم قابل تشخیص نیستند و جاعلان برای این که جعل را حداقل از دید کاربر پنهان کنند، یکسری تبدیلات هندسی مانند: چرخش، تغییر سایز و همچنین دستکاری - هایی مانند فشردن سازی JPEG، نویز گوسین افزایشی و اصلاح گاما را بر روی تصویر اعمال می کنند. بسیاری از روش های تشخیص جعل که تا کنون انجام شده است، نسبت به تمام یا برخی از این تبدیلات مقاوم نیستند.

روش های تشخیص جعل بر پایه ی ویژگی های SIFT نسبت به تمام این تغییرات مقاوم هستند و سرعت محاسباتی بالایی دارند، ولی در صورتی که ناحیه ی کپی شده در قسمت هموار تصویر باشد، این روش ها به خوبی عمل نخواهند کرد.

روش های تشخیص بر پایه ی ممان های زرنیک نسبت به چرخش، و تغییراتی مانند نویز گوسین سفید افزایشی، فشردن سازی JPEG و مات شدگی مقاوم هستند، ولی در مواردی که تغییر سایز و تبدیلات افاین صورت گرفته باشد، به خوبی عمل نمی کنند. در ضمن این روش ها قادر به شناسایی جعل در نواحی هموار تصویر نیز هستند. بنابراین، روشی ترکیبی که از ویژگی های SIFT و ممان های زرنیک برای تشخیص جعل استفاده کند، روشی دقیق خواهد بود که در برابر تغییر سایز، چرخش، نویز گوسین سفید افزایشی، فشردن سازی JPEG، مات شدگی، اصلاح گاما و تبدیلات افاین مقاوم است. از طرفی برای بالا بردن دقت روش پیشنهادی، بهبودهایی روی آن انجام گرفته است که به شرح زیر می باشد:

✓ تشخیص چندین ناحیه ی کپی شده از یک ناحیه. در حالی که روش های قبلی، قادر به تشخیص فقط یک ناحیه ی کپی شده از ناحیه ی اصلی هستند.

✓ این روش بین اشیا و نواحی شبیه به هم، که منتج به استخراج ویژگی های مشابه می شود، و نواحی یا اشیا کپی شده تمییز قائل می شود. با ترکیب روش های تشخیص جعل بر پایه ی ویژگی های SIFT و ممان های زرنیک برای بار اول، توانستیم روش پیشنهادی را علاوه بر مزیت های

گفته شده و بهبودهای انجام شده، در مقابل جعل در نواحی هموار نیز، مقاوم سازیم. همچنین در صورت تشخیص جعل، به تعیین پارامترهای تبدیلات هندسی خواهیم پرداخت. این روش تشخیص جعل دارای دقت بالایی بوده و بار محاسباتی و در نتیجه زمان محاسبه‌ی پایینی خواهد داشت.

۶-۱- تشخیص جعل کپی - انتقال با استفاده از ویژگی‌های SIFT

همانگونه که گفته شد، روش پیشنهادی بر اساس الگوریتم SIFT به استخراج ویژگی‌هایی می‌پردازد که با استفاده از آن‌ها می‌تواند ناحیه‌ی کپی شده و نوع تبدیل انجام گرفته بر روی ناحیه‌ی کپی شده را تشخیص دهد. در واقع بخش کپی شده از لحاظ ظاهری شبیه به بخش چسبانده شده است، بنابراین، نقاط کلیدی که در ناحیه‌ی جعل شده و اصلی به دست می‌آیند کاملاً شبیه به هم هستند. بنابراین، تطبیق ویژگی‌های SIFT می‌تواند جعل احتمالی را تشخیص دهد. این روش به طور خلاصه شامل مراحل زیر است:

مرحله‌ی اول شامل استخراج ویژگی‌های SIFT و تطبیق نقاط کلیدی، مرحله‌ی دوم خوشه‌بندی نقاط کلیدی و تشخیص جعل و مرحله‌ی سوم تخمین تبدیلات هندسی، در صورت تشخیص جعل است.

۶-۱-۱- انطباق نقاط کلیدی

فرض کنید تصویری آزمایشی داده شده است و مجموعه‌ای از نقاط کلیدی $X = \{x_1, \dots, x_n\}$ و توصیف‌کننده‌های SIFT متناظر آن‌ها $\{f_1, \dots, f_n\}$ استخراج شده است. برای تعیین نواحی کپی شده، عملیات انطباق را بین بردارهای f_i انجام می‌دهیم. بهترین کاندید به عنوان نقطه‌ی منطبق با نقطه‌ی کلیدی x_i با تعیین نزدیک‌ترین همسایه‌ی آن از بین $(n-1)$ نقطه‌ی کلیدی دیگر پیدا می‌شود. برای تصمیم‌گیری در مورد این که آیا دو نقطه‌ی کلیدی منطبق هستند یا

نه (یعنی آیا این دو توصیف کننده شبیه به هم هستند یا نه؟)، استفاده از یک حد آستانه‌ی عمومی برای محاسبه‌ی فاصله‌ی ساده‌ی بین دو توصیف‌گر (به عنوان مثال فاصله‌ی اقلیدسی)، روش مناسبی نیست و این به دلیل بعد بالای فضای ویژگی (۱۲۸ تایی) است که در آن برخی ویژگی‌ها نسبت به بقیه، قدرت تمایز بخشی بیشتری دارند.

برای حل این مشکل، در [۷۲] از نسبت فاصله‌ی نزدیک‌ترین همسایه با دومین نزدیک‌ترین همسایه استفاده می‌شود و این نسبت را با یک حد آستانه‌ی T مقایسه می‌کنند (اغلب ۰,۶). برای اثبات این موضوع، فرض کنید که برای نقطه‌ی کلیدی داده شده، بردار D را به صورت زیر تعریف می‌کنیم:

$D = \{d_1, d_2, \dots, d_{n-1}\}$. این بردار، فاصله‌های اقلیدسی بردار توصیف کننده‌ی نقطه‌ی کلیدی را با سایر توصیف کننده‌ها، به صورت مرتب شده نشان می‌دهد. بر طبق این نظریه، یک نقطه‌ی کلیدی در صورتی تطبیق داده می‌شود که رابطه‌ی زیر برقرار باشد:

$$d_1 / d_2 < T \text{ where } T \in (0, 1). \quad (1-6)$$

به این فرآیند، آزمایش 2NN گفته می‌شود.

این فرآیند تطبیق یک مشکل اساسی دارد و آن این است که در مواردی که چندین ناحیه‌ی کپی شده از یک ناحیه وجود داشته باشد، این روش، قادر به شناسایی تمام نواحی کپی شده نخواهد بود. از این رو در این پایان نامه، روش انطباق جدیدی را ارائه می‌دهیم که تغییر یافته‌ی فرمول (۱-۶) است و می‌تواند چندین کپی از ویژگی‌های مشابه را پیدا کند. در تست 2NN، در صورتی که تطبیق صورت گیرد، نسبت فاصله‌ی نزدیک‌ترین همسایه‌ی نقطه‌ی کلیدی به دومین نزدیک‌ترین همسایه، کمتر از ۰,۶، و در موردی که دو ویژگی به صورت تصادفی انتخاب شوند، بزرگ‌تر از ۰,۶ است.

روش ما برای تغییر دادن فرمول (۶-۱) به این صورت است: تست 2NN را آن قدر بین d_i/d_{i+1} تکرار می‌کنیم تا این مقدار بزرگ‌تر از حد آستانه‌ی T شود (ما در آزمایش‌هایمان این مقدار را برابر با ۰.۵ قرار دادیم). اگر k مقداری باشد که در آن فرآیند می‌ایستد، هر نقطه‌ی کلیدی متناظر با یک فاصله در $\{d_1, \dots, d_k\}$ (که $1 \leq k \leq n$) به عنوان تطبیقی برای نقطه‌ی کلیدی مورد آزمایش در نظر گرفته می‌شود.

با تکرار این فرآیند بر روی نقاط کلیدی X مجموعه‌ای از نقاط تطبیق یافته را به دست می‌آوریم. تمام نقاط تطبیق یافته نگه داشته می‌شوند و نقاطی که جدا شده‌اند، در مراحل بعدی در نظر گرفته نمی‌شوند. تا این جا می‌توانیم به صورت نه چندان قطعی در مورد صحت تصاویر نظر بدهیم. اما در مواردی که، تصویر دارای نواحی با بافت مشابه است و نواحی و اشیاء شبیه به هم و کپی نشده در تصویر وجود دارند، ممکن است نقاط کلیدی در این نواحی به اشتباه تطبیق داده شوند. برای رفع این مشکل، دو مرحله‌ی زیر پیشنهاد شده است.

۶-۱-۲- خوشه بندی و تشخیص جعل

در [۷۸] برای تشخیص نواحی کپی شده‌ی محتمل، یک روش خوشه بندی سلسله مراتبی پایین به بالا^۱ بر روی نقاط تطبیق یافته اجرا می‌شود. خوشه بندی سلسله مراتبی، سلسله مراتبی از خوشه‌ها ایجاد می‌کند که توسط یک ساختار درختی نشان داده می‌شود. الگوریتم، با اختصاص دادن هر نقطه‌ی کلیدی به یک خوشه شروع می‌شود، سپس فواصل دو به دویی تمامی خوشه‌ها را در محتوای مکانی محاسبه و نزدیک‌ترین فاصله بین دو خوشه را پیدا می‌کند و در نهایت آن‌ها را در یک خوشه ادغام می‌کند. این عملیات تا آنجا ادامه می‌یابد که به وضعیت پایانی فرآیند ادغام برسیم. این

^۱ Agglomerative Clustering

وضعیت پایانی به دو چیز بستگی دارد: روش ادغام و حد آستانه‌ای که برای گروه بندی خوشه‌ها استفاده می‌شود [۷۸].

روش‌های ادغام مختلفی وجود دارد که در این قسمت به ارزیابی عملکرد هرکدام از آن‌ها پرداخته‌ایم تا بهترین حد آستانه (T) را برای تشخیص جعل تخمین بزنیم. در این قسمت به طور مختصر به بررسی سه روش مختلف ادغام می‌پردازیم: Single، Centroid و Ward. فرض کنید دو خوشه‌ی P و Q وجود دارد که به ترتیب شامل n_P و n_Q عنصر هستند (که در آن x_{Pi} و x_{Qj} به ترتیب به i امین و j امین شیء در خوشه‌های P و Q اشاره می‌کنند)، روش‌های ادغام به صورت‌های زیر عمل می‌کنند [۷۸]:

• روش Single از کوچک‌ترین فاصله‌ی اقلیدسی بین اشیاء در دو خوشه استفاده می‌کند:

$$\text{dist}(P, Q) = \min(\|x_{Pi}, x_{Qj}\|_2) \quad (2-6)$$

with $i = [1, n_P], j = [1, n_Q]$.

• روش Centroid از فاصله‌ی اقلیدسی بین مراکز دو خوشه استفاده می‌کند:

$$\text{dist}(P, Q) = \|\bar{x}_P - \bar{x}_Q\|_2 \quad (3-6)$$

که در آن

$$\bar{x}_Q = \frac{1}{n_Q} \sum_{i=1}^{n_Q} x_{Qi}, \quad \bar{x}_P = \frac{1}{n_P} \sum_{i=1}^{n_P} x_{Pi}, \quad (4-6)$$

• روش Ward افزایش و کاهش مجموع مربعات خطا^۱ را بعد از ادغام کردن دو خوشه در یکی

و با توجه به دو خوشه‌ی جدا محاسبه می‌کند:

$$\Delta_{\text{dist}}(P, Q) = \text{ESS}(PQ) - [\text{ESS}(P) + \text{ESS}(Q)] \quad (5-6)$$

¹ Error Sum of Squares (ESS)

که در آن

$$ESS(P) = \sum_{i=1}^{n_p} \left| x_{Pi} - \bar{x}_P \right|^2 \quad (6-6)$$

x_p مرکز و PQ به خوشه‌ی ترکیب شده اشاره می‌کند.

در این تحقیق، از روش ادغام ward استفاده خواهیم کرد. اکنون باید به دنبال حد آستانه‌ای (T) مناسب برای تعیین تعداد نهایی خوشه‌ها باشیم. نام پارامتری را که با T مقایسه می‌شود ضریب ناسازگاری گذاشته‌ایم که برای تعیین هر عملیات خوشه بندی استفاده می‌شود. هرچقدر مقدار این ضریب بیشتر باشد، اشیایی که توسط لینک به هم متصل می‌شوند، شباهت کمتری با هم دارند، بنابراین، وقتی که مقدار آن از حد آستانه‌ی T بیشتر می‌شود، عملیات خوشه بندی متوقف می‌شود. ضریب ناسازگاری، میانگین فاصله‌ی بین خوشه‌ها را در نظر می‌گیرد و اجازه نمی‌دهد که خوشه‌هایی که در آن سطح از سلسله مراتب، در حوزه‌ی مکانی بسیار از هم فاصله دارند، به هم متصل شوند. تعیین مقدار مناسب برای T به طور مستقیم بر عملکرد تشخیص جعل اثر می‌گذارد [۷۸].

در انتهای فرآیند خوشه بندی، خوشه‌هایی که شامل تعداد مشخصی (سه یا بیشتر از سه تا)، از نقاط کلیدی انطباق یافته نیستند حذف می‌شوند. بر این اساس، برای بهبود عملکرد تشخیص، این نکته را در نظر می‌گیریم که یک تصویر در صورتی دچار جعل کپی- انتقال شده است که دو یا تعداد بیشتری از خوشه‌ها با حداقل سه جفت از نقاط کلیدی انطباق یافته که یک خوشه را به دیگری متصل می‌کنند توسط این روش کشف شوند.

اکنون باید به این نکته توجه شود که ممکن است هیچ نقطه‌ی کلیدی انطباق یافته‌ای در تصویر جعلی به دست نیاید، این مسئله زمانی بوجود می‌آید که ویژگی‌های مهم در ناحیه‌ی جعل شده کشف نشوند. به عنوان مثال، وقتی که یک شیء توسط یک قطعه‌ی صاف و هموار پنهان شده است. در این پایان نامه با استفاده از ممان‌های زرنیک به حل این مشکل می‌پردازیم.

۶-۱-۳- تخمین تبدیلات هندسی

وقتی جعلی بودن یک تصویر اثبات می‌شود، می‌توانیم تبدیلات هندسی بین دو ناحیه‌ی جعل شده و کپی شده را تشخیص دهیم [۴]. مختصات نقاط تطبیق یافته را برای دو ناحیه به ترتیب $\tilde{x}_i = (x, y, 1)^T$ و $\tilde{x}'_i = (x', y', 1)^T$ قرار می‌دهیم. ارتباط هندسی آن‌ها می‌تواند توسط یک affine homography تعریف شود، این ماتریس 3×3 با H نشان داده، به صورت زیر تعریف می‌شود:

$$\begin{pmatrix} x' \\ y' \\ 1 \end{pmatrix} = H \begin{pmatrix} x \\ y \\ 1 \end{pmatrix} \quad (۷-۶)$$

این ماتریس را می‌توان با استفاده از حداقل سه نقطه‌ی انطباق یافته محاسبه کرد. در حقیقت H را با استفاده از تخمین حداکثر احتمال homography می‌توانیم تعیین کنیم [۷۹].

این روش، به دنبال H و جفت نقاط $\hat{x}_i$ و $\hat{x}'_i$ ای است که به طور کامل انطباق یافته‌اند و بر طبق معادله‌ی (۸-۶) باعث کاهش مجموع خطا می‌شوند:

$$\sum_i [d(x_i, \hat{x}_i)^2 + d(x'_i, \hat{x}'_i)^2] \quad \text{که در آن} \quad \hat{x}'_i = H \hat{x}_i \quad \forall i. \quad (۸-۶)$$

نقاطی که به اشتباه منطبق شده‌اند، می‌توانند به طور چشمگیری در تخمین homography موثر باشند. به همین دلیل تخمین قبلی را با استفاده از الگوریتم RANSAC^۱ تخمین می‌زنیم [۸۰].

این الگوریتم به صورت تصادفی، مجموعه‌ای (در روش ما سه جفت نقطه) از نقاط انطباق یافته را انتخاب و H را تخمین می‌زند. سپس تمام نقاط باقی مانده با استفاده از H انتقال داده می‌شوند، آن‌گاه فاصله‌ی این نقاط انتقال داده شده با نقاط منطبق متناظرشان مقایسه می‌شود. اگر این مقدار

^۱ RANdom SAmple Consensus

کمتر یا بیشتر از حد آستانه‌ی مشخص β باشد، به ترتیب به عنوان نقاط inliers و outliers در نظر گرفته می‌شوند.

بعد از یک تعداد تکرار مشخص به اندازه‌ی N_iter ، از بین تبدیلات تخمین زده شده، تبدیلی انتخاب می‌شود که با تعداد بیشتری از نقاط inlier در ارتباط است. در آزمایش‌ها ما N_iter ، ۱۰۰۰ و حد آستانه‌ی β ، ۰.۵ در نظر گرفته می‌شود، این اعداد با توجه به این واقعیت است که ما از روش استاندارد نرمال سازی داده‌ها برای تخمین homography استفاده کرده‌ایم.

وقتی affine homography پیدا می‌شود، تبدیلات چرخش و تغییر سایز می‌تواند با استفاده از تجزیه‌ی این ماتریس، و بردار انتقال با در نظر گرفتن مرکز ثقل تعیین شود. در حقیقت H به صورت زیر نشان داده می‌شود [۷۹]:

$$A = \begin{bmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{bmatrix} \text{ که } H = \begin{bmatrix} A & t \\ O^T & 1 \end{bmatrix} \quad (۹-۶)$$

ماتریس A ترکیبی از تبدیلات چرخش و تغییر سایز است. در حقیقت همواره، این ماتریس می‌تواند به شکل زیر تجزیه شود:

$$A = R(\theta)(R(-\Phi))SR(\Phi) \quad (۱۰-۶)$$

در فرمول بالا $R(\theta)$ و $R(\Phi)$ به ترتیب چرخش با زوایای θ و Φ هستند و $S = \text{diag}(s_1, s_2)$ ، ماتریسی قطری برای تبدیل تغییر سایز است. بنابراین، ماتریس A تعریفی از چرخش با زاویه‌ی Φ ، تغییر سایز به ترتیب با فاکتورهای S_1 و S_2 در جهت محورهای X و Y ، چرخش در جهت عکس با زاویه‌ی $-\Phi$ و در نهایت چرخشی دیگر با زاویه‌ی θ می‌باشد. این تجزیه توسط SVD انجام می‌شود. از آنجایی که ماتریس‌های U و V متعامد هستند، ماتریس A می‌تواند به صورت زیر نوشته شود:

$$A = USV^T = (UV^T)(VSV^T) = R(\theta)(R(-\Phi)SR(\Phi)) \quad (11-6)$$

۶-۲- پایگاه داده

در این بخش، با استفاده از پایگاه داده‌ای به نام MICC-F220، روش پیشنهادی را برای تعیین حد آستانه‌ی T مورد بررسی قرار می‌دهیم.

این پایگاه داده شامل تصاویری با محتوای متفاوت است [۸۱]. تعداد تصاویر موجود در آن ۲۲۰ تصویر می‌باشد که ۱۱۰ تا از تصاویر، اصلی و ۱۱۰ تا جعلی هستند. وضوح تصاویر از 722×480 تا 800×600 متفاوت است و سائز ناحیه‌ی جعل شده، به طور متوسط، $1/2$ درصد سائز کل تصویر است.

برای ایجاد ناحیه‌ی جعل شده، قسمتی از تصویر (به شکل مربع یا مستطیل) را به صورت تصادفی (هم از نظر مکان و هم از نظر اندازه) انتخاب کرده، و پس از اعمال دستکاری‌هایی مانند انتقال، چرخش، تغییرسائز (متقارن یا نامتقارن) و یا ترکیبی از آن‌ها بر روی ناحیه‌ی جعل شده، آن را در قسمت دیگری از همان تصویر می‌چسبانیم.

جدول (۶-۱)، ۱۰ حمله از A تا J به عنوان تبدیلات هندسی برای دستکاری‌های ایجاد شده بر روی تصاویر این پایگاه داده را به طور خلاصه بیان می‌کند. برای هر حمله، چرخش θ را به درجه و فاکتور تغییر سائز اعمال شده بر روی ناحیه‌ی جعل شده بر محورهای x و y را به s_x و s_y نشان می‌دهیم (به طور مثال در حمله‌ی G محورهای x و y به اندازه‌ی ۳۰ درصد تغییر سائز داده شده‌اند و هیچ چرخشی اعمال نشده است).

جدول (۶-۱): ۱۰ ترکیب مختلف از حملات اعمال شده بر روی ناحیه‌ی اصلی بر روی تصاویر پایگاه داده‌ی MICC-F220

Attack	θ^0	S_x	S_y
<i>A</i>	0	1	1
<i>B</i>	10	1	1
<i>C</i>	20	1	1
<i>D</i>	30	1	1
<i>E</i>	40	1	1
<i>F</i>	0	1.2	1.2
<i>G</i>	0	1.3	1.3
<i>H</i>	0	1.4	1.2
<i>I</i>	10	1.2	1.2
<i>J</i>	20	1.4	1.2

۶-۳- کشف جعل

در این قسمت، روش پیشنهادی را برای تعیین بهترین حد آستانه‌ی T تحلیل می‌کنیم. برای این کار از آزمایش 4-fold cross-validation استفاده کردیم. بنابراین، از پایگاه داده‌ی MICC-F220 با ۲۲۰ تصویر، ۱۶۵ تصویر که ۳/۴ مجموعه تصاویر است (۸۲ تصویر جعلی و ۸۳ تصویر اصلی) به طور تصادفی برای انجام عملیات آموزش و برای تعیین حد آستانه T استفاده می‌شود و از ۵۵ تصویر باقی مانده (۱/۴ تعداد کل تصاویر) برای تست عملکرد تعیین جعل روش پیشنهادی استفاده می‌شوند. این عملیات ۴ بار تکرار می‌شود و چهار زیر مجموعه متعلق به مجموعه‌ی آموزش (۳ زیر مجموعه) و مجموعه‌ی تست (۱ زیر مجموعه) به صورت چرخشی عوض می‌شوند، و نتیجه‌ی حاصل، میانگین نتایج آن‌هاست. روش پیشنهادی با استفاده از TPR^1 و FPR^2 ارزیابی می‌شود. TPR نسبت تصاویر جعلی درست تشخیص داده شده به تعداد تصاویر جعلی و FPR نسبت تصاویر اصلی به اشتباه جعلی تشخیص داده شده به تعداد تصاویر اصلی است.

¹ True Positive Rate

² False Positive Rate

$$TPR = \frac{\text{تعداد تصاویر جعلی که به درستی تشخیص داده شده‌اند}}{\text{تعداد تصاویر جعلی}}$$

$$FPR = \frac{\text{تعداد تصاویر اصلی که به اشتباه جعلی داده شده‌اند}}{\text{تعداد تصاویر اصلی}}$$

همان‌طور که قبلاً گفته شد، یک تصویر در صورتی جعلی تشخیص داده می‌شود که توسط این روش، دو (یا تعداد بیشتری) خوشه با حداقل سه جفت از نقاط تطبیق داده شده که خوشه‌ای را به خوشه‌ی دیگر متصل می‌کنند کشف شود.

جدول (۲-۶): فاز آموزش: مقادیر TPR و FPR با توجه به T

T	FPR(%)	TPR(%)
0.8	0.901	10.906
1	4.504	31.820
1.2	6.122	78.937
1.4	7.56	91.955
1.6	8.147	96.368
1.8	8.509	98.114
2	9.560	100
2.2	5.670	100
2.4	7.780	95.89
2.6	7.559	87.978
2.8	3.765	75.174
3	3.789	60.512

در جدول (۲-۶) مقادیر TPR و FPR که با توجه به حد آستانه و در مرحله‌ی آموزش به‌دست آمده است، آورده شده است. این مقادیر در بازه‌ی [۰/۸ و ۳] و با فواصل ۰/۲ هستند. هدف، کاهش FPR و افزایش TPR است. با توجه به جدول مشخص است که FPR همواره کوچک است در حالی - که TPR در مقادیر بزرگ در تغییر است. مقدار بهینه‌ی T، در مقادیر ماکزیمم TPR است که با توجه به جدول مشخص است که این مقدار برای روش ما ۲/۲ است. و در نهایت در مرحله‌ی تست، از مقدار

بهینه‌ی T به دست آمده در مرحله‌ی آموزش استفاده می‌شود. با توجه به آزمایش‌ها انجام شده در مرحله‌ی تست، مقدار FPR برابر با ۵٪ و TPR برابر با ۹۸٪ به دست آمده است. این مقادیر نشان می‌دهد که روش پیشنهادی عملکرد رضایت بخشی دارد به این دلیل که مقدار FPR برای این روش، پایین است و تشخیص جعل را به خوبی انجام می‌دهد.

علاوه بر این، در مواردی که تشخیص جعل به درستی انجام می‌شود به تخمین پارامترهای تبدیلات هندسی می‌پردازیم. برای تعیین میزان دقت تخمین پارامترها از معیار خطای مطلق متوسط^۱ بین پارامترهای تخمین زده شده و مقادیر واقعی آن‌ها استفاده شده است (جدول (۳-۶)). میزان انتقال و تغییر سایز با پیکسل و چرخش با درجه نشان داده شده است.

جدول (۳-۶): خطای تخمین پارامترهای هندسی بر روی تصاویر پایگاه داده‌ی MICC-F220 ($T=2/2$).

نوع تبدیل	انتقال در جهت محور X	انتقال در جهت محور Y	چرخش	فاکتور تغییر سایز در راستای محور X	فاکتور تغییر سایز در راستای محور Y
مقدار MAE	4.04	2.48	0.94	0.021	0.015

نتایج نشان می‌دهد که تخمین پارامترهای تبدیل با دقت بالایی انجام می‌شود.

۴-۶- نتایج به دست آمده

در این بخش به بررسی نتایج حاصل از اجرای روش پیشنهادی بر روی تصاویر مختلف می‌پردازیم:

در تصویر زیر همان‌طور که مشخص است، شخص سمت چپ در تصویر اصلی، به وسیله‌ی قسمتی از پشت زمینه پوشانده شده است که روش پیشنهادی توانسته است این جعل را تشخیص دهد.

^۱ Mean Absolute Error (MAE)

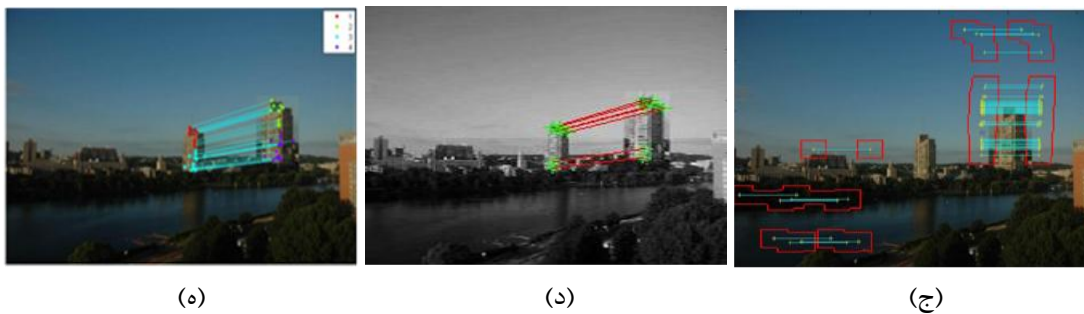


شکل (۶-۱): نمونه‌ای از تشخیص جعل توسط روش پیشنهادی. (الف) تصویر اصلی، (ب) تصویر جعلی، (ج) جعل تشخیص داده شده توسط روش پیشنهادی

۶-۵- مقایسه‌ی روش پیشنهادی با سایر روش‌ها

اکنون روش پیشنهادی را با روش‌های تشخیص جعل بر مبنای ضرایب DCT [۸۲] و ویژگی‌های SIFT [۴] مقایسه می‌کنیم. نتایج حاصله در ادامه آمده است:

تغییر سایز: در شکل (۶-۲)، برج موجود در تصویر کپی شده، تغییر سایز داده شده و در کنار برج قبلی چسبانده شده.



شکل (۶-۲): تغییر سایز در ناحیه‌ی جعل شده. (الف) تصویر اصلی، (ب) تصویر جعل شده (تغییر سایز در ناحیه‌ی کپی شده)، (ج) تشخیص جعل توسط (ج) روش DCT [۸۲]، (د) روش SIFT [۴]، و (ه) روش پیشنهادی

همان‌طور که مشخص است روش DCT [۸۲] نتوانسته این جعل را به درستی تشخیص بدهد، بنابراین، این روش در برابر تغییر سایز مقاوم نیست. روش SIFT [۴] با تعداد نقاط تطبیق یافته‌ی کمتر و روش پیشنهادی، به درستی و با تعداد نقاط تطبیق یافته‌ی بیشتری، این جعل را تشخیص داده است.

چرخش: در تصویر زیر، تصویر شیء مکعبی چوبی، کپی شده، به اندازه‌ی ۳۰ درجه به سمت منفی محور X ها چرخیده و در سمت راست تصویر چسبانده شده است.



(ب)

(الف)



(ه)



(د)



(ج)

شکل (۳-۶): چرخش در ناحیه‌ی جعل شده. الف) تصویر اصلی، ب) تصویر جعل شده (چرخش در ناحیه‌ی کپی شده)، تشخیص جعل توسط ج) روش DCT [۸۲]، د) روش SIFT [۴]، و ه) روش پیشنهادی

همان‌طور که مشخص است روش DCT [۸۲] نتوانسته تشخیص درست بدهد، بنابراین، این روش در برابر چرخش نیز مقاوم نیست؛ اما روش SIFT [۴] و روش پیشنهادی به خوبی توانستند این جعل را تشخیص دهند.

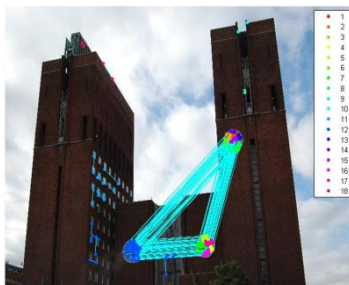
ایجاد چندین ناحیه‌ی کپی شده از یک ناحیه: در شکل زیر، از تصویر ساعت، سه کپی ایجاد و در قسمت‌های مختلف چسبانده شده است.



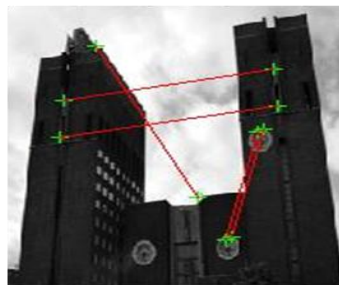
(ب)



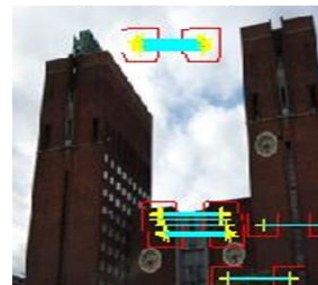
(الف)



(ه)



(د)



(ج)

شکل (۴-۶): چند ناحیه‌ی کپی شده از یک ناحیه. (الف) تصویر اصلی، (ب) تصویر جعل شده (ایجاد چندین شیء کپی شده از یک شیء)، تشخیص جعل توسط (ج) روش DCT [۸۲]، (د) روش SIFT [۴]، و (ه) روش پیشنهادی

همان‌طور که مشخص است روش DCT [۸۲] و روش SIFT [۴] نتوانسته‌اند هر سه شیء کپی شده را تشخیص دهند، بنابراین، این روش‌ها در مواردی که چندین کپی از یک ناحیه یا شیء ایجاد شده باشند، به درستی عمل نخواهند کرد، با این حال روش پیشنهادی به خوبی هر سه شیء کپی شده را تشخیص داده است.

نویز گوسین سفید افزایی: در تصویر زیر، یکی از بوته‌های گل، کپی شده و در سمت

راست تصویر چسبانده شده است و سپس نویز گوسین سفید با $dB=20$ به تصویر افزوده شده است.



(ب)



(الف)



(د)



(ج)

شکل (۵-۶): (الف) افزودن نویز گوسین سفید. تصویر اصلی، (ب) تصویر جعل شده، (ج) تصویر حاصل از اعمال نویز گوسین سفید افزایشی با $dB=20$ ، (د) تشخیص جعل توسط روش پیشنهادی.

همان‌طور که مشخص است، روش پیشنهادی نسبت به نویز گوسین سفید مقاوم است.

فشرده سازی JPEG: در تصویر زیر، مارک روی یکی از فنجان‌ها کپی شده و بر روی فنجان

دیگر چسبانده شده است. آن‌گاه تصویر را با استفاده از فشرده سازی JPEG و با فاکتور کیفی $Q=40$

فشرده کردیم.



(ب)



(الف)



(د)



(ج)

شکل (۶-۶): فشرده سازی JPEG تصویر. الف) تصویر اصلی، ب) تصویر جعل شده، ج) تصویر حاصل از اعمال فشرده سازی JPEG با فاکتور کیفی $Q=40$ ، د) تشخیص جعل توسط روش پیشنهادی.

همان طور که مشخص است، روش پیشنهادی نسبت به فشرده سازی JPEG مقاوم است.

گوسین بلورینگ: در تصویر زیر، یکی از تفنگ ها کپی شده، و در قسمت دیگر چسبانده شده

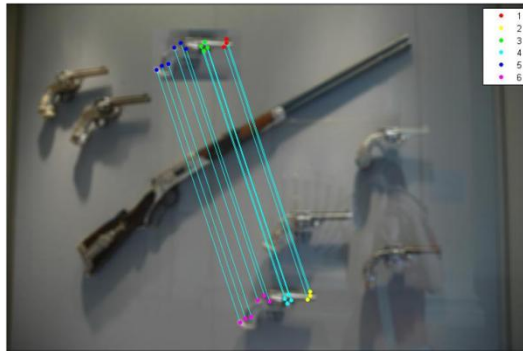
است، سپس تصویر حاصل توسط گوسین بلورینگ با $\sigma=7$ و $\sigma=35$ دستکاری شده است.



(ب)



(الف)



(د)



(ج)

شکل (۶-۷): مات کردن تصویر. (الف) تصویر اصلی، (ب) تصویر جعل شده، (ج) تصویر حاصل از اعمال گوسین بلورینگ با $\omega=7$ و $\sigma=35$ ، (د) تشخیص جعل توسط روش پیشنهادی.

همان طور که مشخص است، روش پیشنهادی نسبت به گوسین بلورینگ مقاوم است.

اصلاح گاما: در تصویر زیر، یک کپی از قطعه سنگ ایجاد و بر روی چمن چسبانده ایم و سپس

تصویر را با استفاده از روش اصلاح گاما با $\gamma=0.1$ ، بهبود داده ایم.



(ب)



(الف)



(د)



(ج)

شکل (۸-۶): بهبود توسط اصلاح گاما. (الف) تصویر اصلی، (ب) تصویر جعل شده، (ج) تصویر بهبود یافته توسط اصلاح گاما با $\gamma=0.1$ ، (د) تشخیص جعل توسط روش پیشنهادی.

همان‌طور که مشخص است این روش در برابر اصلاح گاما نیز مقاوم است.

تبدیل افاین: تبدیل افاین، تبدیلی است که در آن، طول و عرض یک شیء به نسبت غیر

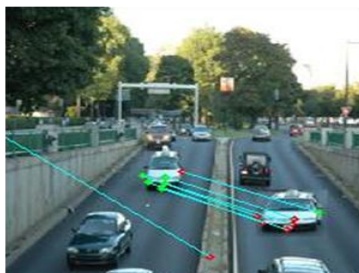
مساوی، بزرگ‌تر یا کوچک‌تر می‌شوند. به طور مثال، طول ۲ برابر و عرض ۴ برابر و یا این‌که بر روی

یک ناحیه و یا یک شیء ترکیبی از تبدیلات مختلف مانند تغییر سایز، چرخش و انتقال انجام شود.

در شکل زیر، تصویر ماشین، کپی شده، سپس تبدیل افاین بر روی آن انجام شده (فاکتور تغییر

سایز اعمال شده بر روی آن در راستای محور X دو برابر محور Y است) و در قسمت دیگری از تصویر

چسبانده شده است.



(الف)



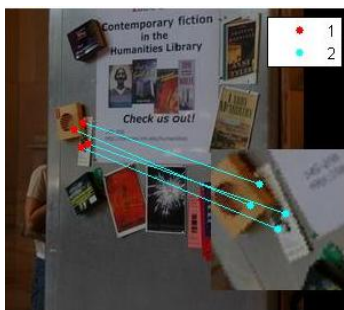
(ب)



(ج)

شکل (۶-۹): تبدیل افاین در ناحیه‌ی جعل شده. الف) تصویر اصلی، ب) تصویر جعل شده (اعمال تبدیل افاین بر روی ناحیه‌ی کپی شده)، د) تشخیص جعل توسط روش پیشنهادی.

در شکل زیر، تصویر مکعب چوبی، کپی شده، سپس تبدیل افاین بر روی آن انجام شده (چرخش به اندازه‌ی ۳۰ درجه و شیء در راستای طول و عرض دو برابر شده است) و در قسمت دیگری از تصویر چسبانده شده است.



(الف)



(ب)



(ج)

شکل (۶-۱۰): تبدیل افاین در ناحیه‌ی جعل شده. الف) تصویر اصلی، ب) تصویر جعل شده (اعمال تبدیل افاین بر روی ناحیه‌ی کپی شده)، د) تشخیص جعل توسط روش پیشنهادی.

همان‌طور که مشخص است، این روش در برابر تبدیلات افاین نیز مقاوم است.

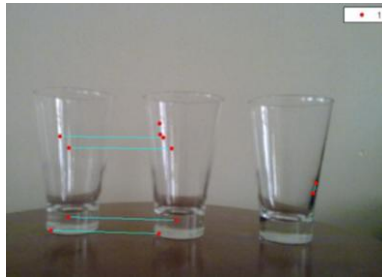
نواحی شبیه به هم: در تصویر زیر چند شیء مشابه وجود دارد. از آنجایی که این اشیاء از

لحاظ ظاهری بسیار به هم شبیه هستند، منتج به استخراج ویژگی‌های SIFT یکسان می‌شوند.

این تصویر با استفاده از دو روش [۴] و روش پیشنهادی آزمایش شده است.



(الف)



(ج)



(ب)

شکل (۶-۱۱): نواحی شبیه به هم. الف) تصویر اصلی (شامل اشیاء شبیه به هم، و کپی نشده)، ب) نتایج حاصل از روش SIFT [۴]، ج) نتایج حاصل از روش پیشنهادی.

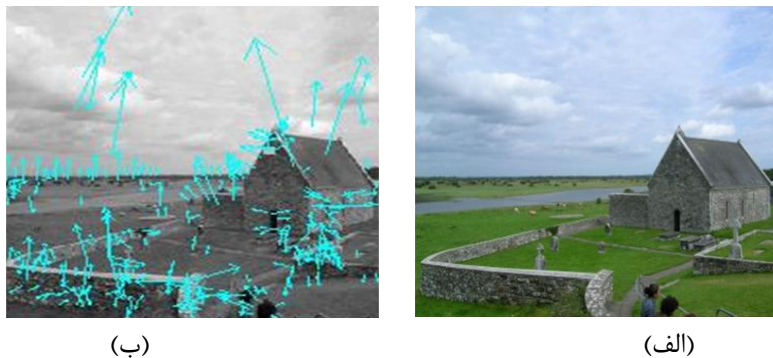
همان‌طور که از نتایج مشخص است، روش [۴]، تصویر را به اشتباه، جعلی تشخیص داده در حالی که روش پیشنهادی، صحت تصویر را اعلام کرده است. (روش پیشنهادی در صورتی تصویری را جعلی اعلام می‌کند که حداقل دو خوشه با حداقل سه جفت نقطه‌ی کلیدی تطبیق داده شده، در هر خوشه در تصویر یافت شود).

با توجه به نتایج، مشخص است که روش پیشنهادی در برابر تغییر سایز، چرخش، نویز گوسین سفید افزایشی، فشرده سازی JPEG، مات شدگی، اصلاح گاما و تبدیل افاین مقاوم است. همچنین در مواردی که چندین ناحیه‌ی کپی شده از یک ناحیه ایجاد می‌شود، به درستی عمل خواهد کرد.

نواحی هموار: همان‌طور که قبلاً هم اشاره شد، از آنجایی که در قسمت‌های هموار تصویر، هیچ نقطه‌ی کلیدی (مانند SIFT) به دست نمی‌آید، روش‌های تشخیص جعل با استفاده از نقاط کلیدی، از

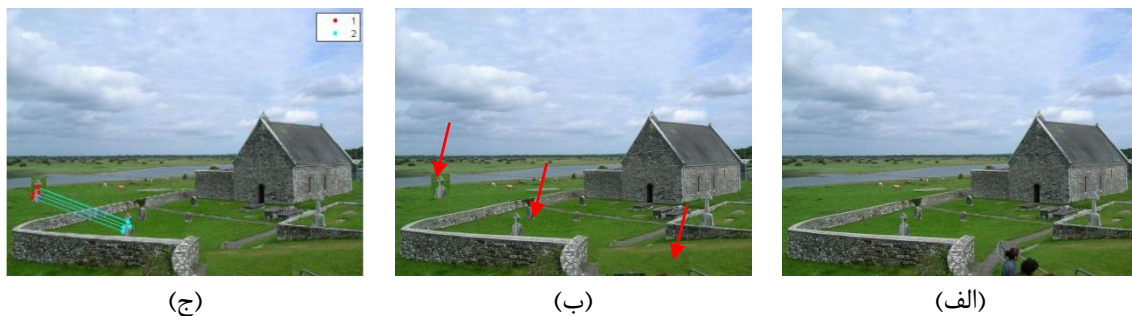
جمله روش‌های بر پایه‌ی ویژگی‌های SIFT، قادر به شناسایی نواحی کپی شده در قسمت‌های هموار تصویر نخواهند بود.

با دقت در تصویر زیر متوجه خواهیم شد که در نواحی هموار مانند چمن و آسمان، نقطه‌ی کلیدی (SIFT) به‌دست نیامده است:



شکل (۶-۱۲): استخراج ویژگی‌های SIFT. الف) تصویری که شامل نواحی هموار مانند چمن و آسمان است، ب) نقاط SIFT به‌دست آمده.

اکنون قسمتی از ناحیه‌ی هموار (چمن) را کپی کرده و به منظور پنهان کردن، بر روی تصویر دو شخص موجود در پایین تصویر قرار می‌دهیم، همچنین شیء موجود در قسمت پایین را کپی کرده و در سمت چپ تصویر قرار می‌دهیم.



شکل (۶-۱۳): تشخیص جعل با استفاده از فقط ویژگی‌های SIFT. الف) تصویر اصلی، ب) تصویر جعل شده، د) جعل تشخیص داده شده توسط روش پیشنهادی (با استفاده از فقط ویژگی‌های SIFT)

همان‌طور که مشخص است روش تشخیص جعل تنها با استفاده از ویژگی‌های SIFT، جعل چمن را (به دلیل هموار بودن بافت) تشخیص نداده، ولی جعل شیء را توانسته به خوبی تشخیص دهد. پس باید به دنبال راه حل دیگری باشیم.

۶-۶- تشخیص جعل کپی - انتقال با استفاده از ممان‌های زرنیک [۳]

یکی از ویژگی‌های ممان‌های زرنیک که در این پایان نامه برای ما حائز اهمیت است، وجود آن‌ها در نواحی هموار تصویر می‌باشد؛ بنابراین، با استفاده از آن‌ها می‌توانیم به تشخیص جعل در نواحی هموار تصویر بپردازیم.

برای تشخیص جعل، این روش باید معادله‌ی (۵-۹) را از دیدگاه جبری بر آورده کند. علاوه بر این، از آنجایی که اغلب جاعلان برای پنهان کردن جعل خود، تصاویر را دستکاری می‌کنند، این روش باید نسبت به نویز افزایشی یا مات شدگی مقاوم باشد که در [۱۴] آورده شده است. ابتدا تصویر $M \times N$ را برای محاسبه‌ی ممان‌های زرنیک، به زیر بخش‌های هم پوشانی $L \times L$ تقسیم می‌کنیم. هر بلوک به صورت B_{ij} نشان داده می‌شود. که در آن i و j به ترتیب به سطر و ستون شروع بلوک اشاره می‌کنند.

$$B(x, y) = f(x + i, y + j) \quad (۱۲-۶)$$

$$j \in \{0, \dots, N - 1\}$$

$$x, y \in \{0, \dots, L - 1\}, i \in \{0, \dots, M - 1\},$$

بنابراین، می‌توانیم N_{blocks} زیر بخش‌های هم پوشانی را به دست آوریم:

$$N_{\text{blocks}} = (M - L + 1) \times (N - L + 1) \quad (۱۳-۶)$$

فرض می‌کنیم که سائز از پیش تعریف شده‌ی بلوک، کوچک‌تر از ناحیه‌ی جعل شده باشد.

بنابراین، ممان‌های زرنیک درجه‌ی n از بلوک محاسبه می‌شوند و توسط تابع `getZernikeMoments` به صورت زیر به شکل برداری در می‌آیند:

$$V_{ij} = \text{getZernikeMoments}(B_{ij}, n), \quad (14-6)$$

که

function `V=getZernikeMoments (Block,nMax)`

```

1: Vector V
2: for n=0 to n Max do
3:   for m=0 to n do
4:     if (n-m)%2=0 then
5:       V.pushback ( $\frac{n+1}{\pi} \sum \sum (Block \times V_{nm}^*)$ )
6:     end if
7:   end for
8: end for
9: return V

```

با تحلیل تابع `getZernikeMoments` از مرتبه‌ی داده شده‌ی $nMax$ ، تعداد تمام مومنت‌ها

برابرست با:

$$N_{moments} = \sum_{i=0}^{nMax} \left(\left\lfloor \frac{i}{2} \right\rfloor + 1 \right) \quad (15-6)$$

بعد از آن مجموعه‌ای از ممان‌های زرنیک برداری شده را در ماتریسی به نام Z به صورت زیر

نشان می‌دهیم:

$$Z = \begin{bmatrix} V_{00} \\ \dots \\ V_{(M-L)(N-L)} \end{bmatrix} \quad (16-6)$$

سپس، ماتریس Z با استفاده از روش مرتب سازی بر مبنای واژه نگاری مرتب می‌شود. این

مجموعه‌ی مرتب شده را با $\hat{Z}$ نشان می‌دهیم. در این مجموعه، فاصله‌ی اقلیدسی بین هر جفت بردار

همسایه را محاسبه می‌کنیم. اگر این فاصله از کوچک‌تر از حد آستانه‌ی از پیش تعریف شده‌ی D_1 باشد، آن بردارهای همسایه را کاندیدی برای جعل در نظر می‌گیریم:

$$\hat{Z}_p = (\hat{Z}_1^p, \hat{Z}_2^p, \dots, \hat{Z}_{N_{moments}-1}^p, \hat{Z}_{N_{moments}}^p),$$

$$\hat{Z}_{p+1} = (\hat{Z}_1^{p+1}, \hat{Z}_2^{p+1}, \dots, \hat{Z}_{N_{moments}-1}^{p+1}, \hat{Z}_{N_{moments}}^{p+1}), \quad (17-6)$$

$$\sqrt{\sum_{q=1}^{N_{moments}} (z_q^p - z_q^{p+1})^2} < D_1$$

با توجه به این موضوع که بلوک‌های همسایه دارای ممان‌های زرنیک مشابه هستند، فاصله‌ی بین بلوک‌های واقعی تصویر را به صورت زیر محاسبه می‌کنیم:

$$\sqrt{(i-k)^2 + (j-1)^2} > D_2 \quad (18-6)$$

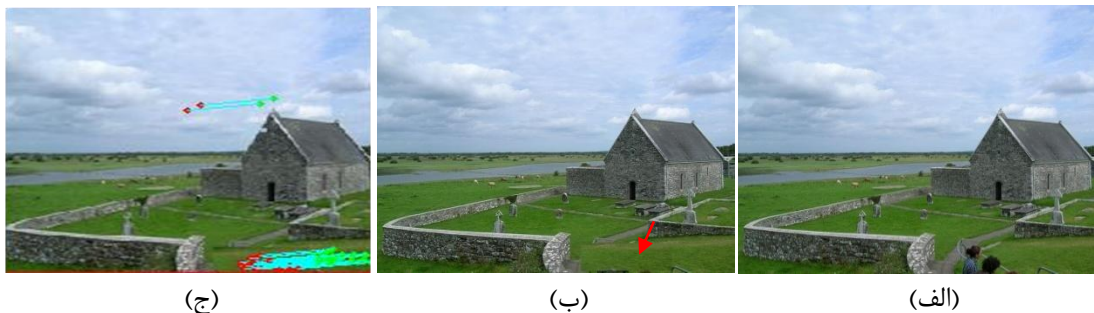
$$\hat{Z}_p = V_{ij} \text{ and } \hat{Z}_{p+1} = V_{kl}$$

در نهایت با استفاده از معادلات (17-6) و (18-6) در مورد این که آیا بلوک‌های بررسی شده جعلی هستند یا نه تصمیم می‌گیریم.

۶-۷- نتیجه‌ی حاصل از ترکیب ویژگی‌های SIFT و ممان‌های زرنیک

در ادامه، نتیجه‌ی حاصل از اجرای روش پیشنهادی که تشخیص جعل با استفاده از ویژگی‌های SIFT و ممان‌های زرنیک است، قرار داده شده است.

در تصویر زیر، قسمتی از چمن را کپی کرده و بر روی دو شخص موجود در قسمت پایین تصویر قرار می‌دهیم.



شکل (۶-۱۴): تشخیص نواحی هموار با استفاده از روش پیشنهادی. الف) تصویر اصلی، ب) تصویر جعل شده، ج) جعل تشخیص داده شده توسط روش پیشنهادی (ترکیب ویژگی‌های SIFT و ممات‌های زرنیک)

با توجه به تصویر مشخص است که این روش، علاوه بر مزیت‌های گفته شده، توانایی تشخیص جعل در نواحی هموار تصویر را نیز دارد.

۶-۸- بررسی میزان مقاومت روش پیشنهادی در برابر فشرده سازی

JPEG، افزودن نویز و اصلاح گاما

برای بررسی این اثرات، آزمایش زیر را ترتیب دادیم. شروع آزمایش با ۱۰ تصویر است. یک بلوک به صورت تصادفی از هر تصویر انتخاب می‌شود (همان‌طور که قبلاً توضیح داده شده)، و بر روی هر کدام ۴ تبدیل هندسی (A، C، D و J از جدول (۶-۲)) اعمال می‌شود. علاوه بر این، قبل از چسباندن آن‌ها، ۲ اصلاح گامای مختلف با مقادیر $[1/4, 2/2]$ به هر بلوک اعمال می‌شود. در نهایت ۸۰ تصویر جعلی به دست می‌آید. در روشی مشابه، مرحله‌ی نهایی اصلاح گاما ابتدا با فشرده سازی JPEG با ۲ فاکتور کیفی مختلف $[50, 75]$ و سپس با افزودن نویز گوسین با SNR (dB) برابر با ۳۰، ۲۰ جایگزین می‌شود. در هر مورد ۸۰ تصویر جعلی به وجود آمده است. بنابراین، برای هر کدام از سه دستکاری (اصلاح گاما، فشرده سازی JPEG، و افزودن نویز)، پایگاه داده‌ای شامل ۸۰ تصویر جعلی و

۱۰۰ تصویر اصلی به صورت تصادفی از پایگاه داده‌ی MICC-F220 ساخته شده است. در جدول

(۴-۶) عملکرد روش پیشنهادی در این آزمایش‌ها به صورت TPR و FPR گزارش شده است.

جدول (۴-۶): ارزیابی عملکرد تشخیص در برابر اصلاح گاما، فشرده سازی JPEG، و افزودن نویز اعمال شده بر روی ناحیه‌ی جعل شده و تبدیل یافته.

عملیات	FPR(%)	TPR(%)
اصلاح گاما	9.01	98.67
فشرده سازی JPEG	10.15	99.12
نویز گوسین سفید	14.56	100.00

این نتایج نشان می‌دهد که روش پیشنهادی، در برابر انواع پس پردازش‌های اعمال شده بر روی

ناحیه‌ی جعل شده که دستخوش تبدیلات هندسی نیز شده است، دقت قابل قبولی دارد.

فصل هفتم

نتیجه‌گیری و پیشنهادها برای پژوهش‌های آینده

۷-۱- نتیجه گیری

امروزه به دلیل در دسترس بودن نرم افزارهای ویرایش تصاویر و سهولت استفاده از آنها، تصاویر دیجیتالی در معرض انواع دستکاری‌ها قرار می‌گیرند. از طرفی، به دلیل وسعت استفاده از تصاویر دیجیتالی مانند روزنامه‌ها، مجله‌ها، اینترنت و ...، این تصاویر به عنوان منبع اطلاعات محسوب می‌شوند، بنابراین، امروزه تشخیص صحت یا عدم صحت تصاویر به عنوان موضوع مهم و اساسی مطرح است.

این موضوع زمانی اهمیت بیشتری پیدا می‌کند که بخواهیم از این تصاویر به عنوان مدرک در دادگاه‌ها استفاده کنیم. بنابراین، برای حفظ حقوق شخصی افراد باید به دنبال روش‌هایی باشیم که اصل و یا جعل بودن تصاویر را مشخص کند. روش‌های مختلفی برای جعل و همچنین تشخیص جعل تصاویر وجود دارد که در این پایان نامه به بررسی روش‌های تشخیص جعل کپی-انتقال پرداختیم.

اغلب جاعلان برای این که جعل خود را از دید کاربر پنهان کنند و هیچ گونه رد پایی بر جای نگذارند، بر روی ناحیه‌ی جعل شده تبدیلات هندسی مانند تغییر سایز، چرخش و همچنین تغییراتی مانند افزودن نویز، اصلاح گاما، فشرده سازی JPEG، مات کردن و... را انجام می‌دهند. بنابراین، برای ارائه‌ی یک روش تشخیص جعل قوی، باید به دنبال روش‌هایی باشیم که در برابر این تبدیلات و تغییرات مقاوم باشند. این کار مستلزم این است که از ویژگی‌هایی استفاده کنیم که در برابر این تغییرات ثابت باشند.

تا کنون روش‌های زیادی برای تشخیص جعل کپی انتقال ارائه شده است که به بررسی اکثر آنها پرداختیم. اغلب روش‌های موجود در برابر تمام یا برخی از این تغییرات مقاوم نیستند و یا زمان محاسبه‌ی بالایی دارند.

در این تحقیق برای ارائه‌ی روش مقاوم در برابر تمام این دستکاری‌ها، روش ترکیبی جدیدی پیشنهاد داده‌ایم که در آن برای تشخیص جعل از ویژگی‌های SIFT و ممان‌های زرنیک استفاده می‌-

کند. ویژگی‌های SIFT در برابر انواع دستکاری‌ها مانند چرخش، تغییر سایز، فشرده سازی JPEG، نویز گوسین افزایشی، اصلاح گاما و تبدیلات آفاین مقاوم است.

ممان‌های زرنیک نسبت به چرخش، و تغییراتی مانند نویز گوسین سفید افزایشی، فشرده سازی JPEG و مات شدگی مقاوم هستند؛ ولی در مواردی که تغییر سایز و تبدیلات آفاین صورت گرفته باشد، به خوبی عمل نمی‌کنند. در ضمن این روش قادر به شناسایی جعل در نواحی هموار تصویر نیز هست.

بنابراین، روش پیشنهادی که از ترکیب دو روش تشخیص جعل بر پایه‌ی ممان‌های زرنیک و ویژگی‌های SIFT حاصل شده، در برابر تغییر سایز، چرخش، نویز گوسین سفید افزایشی، فشرده سازی JPEG، مات شدگی، اصلاح گاما و تبدیلات آفاین مقاوم است. همچنین با توجه به بهبودهایی که بر روی روش پیشنهادی اعمال شده است، می‌تواند چندین ناحیه‌ی کپی شده از یک ناحیه را تشخیص دهد. در حالی‌که روش‌های قبلی، قادر به تشخیص فقط یک ناحیه‌ی کپی شده از ناحیه‌ی اصلی هستند. همچنین این روش بین اشیاء و نواحی شبیه به هم که منتج به ویژگی‌های استخراج شده‌ی مشابه می‌شود و نواحی یا اشیاء کپی شده تمییز قائل می‌شود.

از طرفی، با ترکیب روش‌های تشخیص جعل بر پایه‌ی ویژگی‌های SIFT و ممان‌های زرنیک برای بار اول، توانستیم روش پیشنهادی را علاوه بر مزیت‌های گفته شده و بهبودهای انجام شده، در مقابل جعل در نواحی هموار نیز، مقاوم سازیم.

در این تحقیق، در صورت تشخیص جعل، به تعیین پارامترهای تبدیلات هندسی نیز پرداختیم. این روش تشخیص جعل دارای دقت بالایی بوده، بار محاسباتی و در نتیجه زمان محاسبه‌ی کمی خواهد داشت.

برای اثبات مزایای عنوان شده، به ارزیابی این روش با استفاده از پایگاه داده‌ی MICC-F220 پرداختیم و در نهایت آزمایش‌ها نشان داد که روش پیشنهادی FPR ای برابر با ۵٪ و TPR ای برابر با ۹۷٪ دارد. بنابراین، دقت بالای این روش را به اثبات رساندیم.

۷-۲- پیشنهادها برای پژوهش‌های آینده

برای ادامه‌ی کار در آینده پیشنهادهایی به صورت زیر ارائه کرده‌ایم:

- ✓ ممان‌های زرنیک نسبت به تغییر سایز مقاوم نیستند، بنابراین، می‌توان این روش‌ها را طوری تغییر داد که نسبت به تغییر سایز مقاوم شوند. به عنوان مثال، اگر جعل در ناحیه‌ی هموار صورت گرفته باشد و دچار تغییر سایز هم شده باشد، بتوان آن را تشخیص داد.
- ✓ در این تحقیق از دو ویژگی برای تشخیص جعل استفاده شده است، بنابراین، می‌توان یک ویژگی برای تشخیص جعل انتخاب کرد به طوری که تمام مزایای دو ویژگی را داشته باشد. در این صورت زمان پردازش کاهش خواهد یافت.
- ✓ به دنبال روش‌هایی برای پیدا کردن سایر جعل‌های انجام شده بر روی تصاویر مانند مونتاژ، حذف کردن و ... باشیم، به طوری که تمام حالت‌های ممکن جعل را در نظر بگیرد، به عنوان مثال، روشی که بتواند دو ناحیه‌ی مونتاژ شده از یک ناحیه، مونتاژ در ناحیه‌ی هموار و ... را تشخیص دهد.

1. Fridrich J., Soukal D. and Luk J., **(2003)** "Detection of copy-move forgery in digital images" Proc. Digital Forensic, Research Workshop, Cleveland, OH. 110, 31, pp **44-50**.
2. **Johnson M. K., (2007), Phd Thesis, "Lighting and Optical Tools for Image Forensics" Computer Department, Dartmouth College, Hanover, New Hampshire.**
3. RyuS., LeeM., and LeeH., **(2010)** "Detection of Copy-Rotate-Move Forgery using Zernike Moments," Information Hiding Conference, pp. **51-65**.
4. X. Pan and S. Lyu ,(2010) "Region Duplication Detection Using Image Feature Matching, IEEE Transactions on Information Forensics and Security, vol. 5, no. 4, pp. 857-867, Dec..
5. StampM. And StavroulakisP., **(2010)** "**Handbook of Information and Communication Security**", Springer-Verlag Berlin Heidelberg, pp **810**.
6. HafnerK., **(2004)** "The camera never lies", but the software can", New York Times.
7. MitchellW.J., **(1994)** "When is seeing believing", Sci. Am. 2, 1, pp **44-49**.
8. AmeriniI., BallanL., CaldelliR., BimboA. D., and SerraG., **(2011)** "A SIFT-Based Forensic Method for Copy-Move Attack Detection and Transformation Recovery," IEEE Tansactions on Information Forensics and Security, 6, pp **1099-1110**.
9. Dino A., **(1999)** "Photo fakery: the history and techniques of photographic deception and manipulation". Brassey's, Dulles, VA.
10. Jaubert A., **(1992)**, "Le Commissariat aux Archives". Bernard Barrault.
11. Hwang W. S, **(2004)** "Evidence of a pluripotent human embryonic stem cell line derived from a cloned blastocyst". Science, 303, 5664, pp **1669-1674**.
12. Kennedy D., **(2006)** "Editorial retraction". Science, 311, 5759, pp **335**.
13. Wade N., **(2005)** "Korean scientist said to admit fabrication in a cloning study". The New York Times.
14. Cook G., **(2006)** "Technology seen abetting manipulation of research". The Boston Globe.
15. Liu G, Wang J, Lian S, Dai Y., **(2013)** "Detect image splicing with artificial blurred boundary", 57, 12, pp **2647-2659**.

16. SencarH.T., MemonN., **(2008)** "Overview of State-of-the Art in Digital Image Forensics, Statistical Science and Interdisciplinary Research" (World Scientific Press, Singapore)
17. PopescuA.,and FaridH.,**(2004)** "Statistical tools for digital forensics," 6th International Workshop on Information Hiding, Toronto, Canada.
18. CoxI.J., MillerM.L, BloomJ.A., **(2002)** "Digital Watermarking" Morgan Kaufmann, San Francisco, CA.
19. Cox I.J., Miller M.L., Linnartz J.M.G, Kalker T., **(1999)** "A review of watermarking principles and practices" Digital Signal Processing for Multimedia Systems, ed.byK.K.Parhi, T.Nishitani(MarcelDekker, New York, USA) pp **461–482**.
20. BlytheP., FridrichJ., **(2004)** "Secure digital camera" DigitalForensic Research Workshop, Baltimore, Maryland.
21. LinC.Y., ChangS.F., **(1998)** "A robust image authentication algorithm surviving JPEG lossy compression" Proc. 3312, pp **296–307**.
22. Fridrich J., **(1998)** "Image watermarking for tamper detection" International Conference on Image Processing, 2, pp **404–408**.
23. Kundur, D and Hatzinakos, D., **(1999)** "Digital watermarking for telltale tamper proofing and authentication" Proceedings of the IEEE (USA), 87, 7, pp **1167–1180**.
24. Xiang Zhou, XiaohuiDuan, and Daoxian Wang. **(2004)** "A semi-fragile watermark scheme for image authentication" In Proceedings of the 10th International Multimedia Modelling Conference, page **374**.
25. PopescuA.C., FaridH., **(2005)** "Exposing digital forgeries by detecting traces of re-sampling" IEEE Trans. Signal Process. 53(2), pp **758–767**.
26. DempsterA., LairdN., and RubinD., **(1977)** "Maximum likelihood from incomplete data via the EM algorithm," Journal of the Royal Statistical Society, 99, 1, pp **1–38**.
27. BarniM., CostanzoA., **(2012)** "A fuzzy approach to deal with uncertainty in image forensics", Signal Processing: Image Communication, 27, pp **998–1010**.
28. Devi MahalakshmiS., VijayalakshmiK., PriyadharsiniS., **(2012)** "Digital image forgery detection and estimation by exploring basic image manipulations", Digital Investigation, 8,PP **215–225**.
29. JohnsonM.K., FaridH., **(2005)** "Exposing digital forgeries by detecting inconsistencies in lighting", Proc.ACM Multimedia Security Workshop.
30. SutcuY., CoskunB., SencarH.T., and MemonN., **(2007)** "Tamper detection based on regularity of wavelet transform coefficients" IEEE Int. Conference on Image Processing.

31. J. Redi, W. Taktak, and J.-L. Dugelay, (2011) "Digital Image Forensics: A Booklet for Beginners," *Multimedia Tools and Applications*, 51, 1, pp **133–162**.
32. Muhammad G., Hussain M, Bebis G., (2012) "Passive copy move image forgery detection using undecimated dyadic", *Digital Investigation*, 9, pp **49–57**.
33. Li L., Li Sh., Zhu H, (2013) "An Efficient Scheme for Detecting Copy-move Forged Images by Local Binary Patterns", *Journal of Information Hiding and Multimedia Signal Processing*, 4, 1, pp **789–787**.
34. Christlein V., Riess Ch., Jordan J., Riess C., and Angelopoulou E, (2012) "An Evaluation of Popular Copy-Move Forgery Detection Approaches".
35. BasharM., NodaK., OhnishiN., and MoriK., (2010) "Exploring Duplicated Regions in Natural Images," *IEEE Transactions on Image Processing*, 25, 5, pp **89–95**.
36. BayramS., AvcıbasI., SankurB., and MemonN., (2005) "Image Manipulation Detection with Binary Similarity Measures," in *European Signal Processing Conference*.
37. BayramS., SencarH., and MemonN., (2009) "An efficient and robust method for detecting copy-move forgery," in *IEEE International Conference on Acoustics, Speech, and Signal Processing*, pp **1053–1056**.
38. Bravo-Solorio S., and NandiA. K., (2011) "Exposing Duplicated Regions Affected by Reflection, Rotation and Scaling," in *International Conference on Acoustics, Speech and Signal Processing*, pp **1880–1883**.
39. DybalaB., JenningsB., and LetscherD., (2007) "Detecting Filtered Cloning in Digital Images" in *Workshop on Multimedia & Security*, pp **43–50**.
40. FridrichJ., SoukalD., and LukasJ., (2003) "Detection of copy-move forgery in digital images," in *Proceedings of Digital Forensic Research Workshop*.
41. Lynch G., ShihY., and Liao H.Y., (2013) "An efficient expanding block algorithm for image copy-move forgery detection", *Information Sciences*, 239, pp **253–265**.
42. Ju S., Zhou J., and He K., (2007) "An Authentication Method for Copy Areas of Images," *International Conference on Image and Graphics*, pp **303–306**.
43. Kang X. and WeiS., (2008) "Identifying Tampered Regions Using Singular Value Decomposition in Digital Image Forensics," *International Conference on Computer Science and Software Engineering*, 3, pp **926–930**.
44. KeY., SukthankarR., and HustonL., (2004) "An Efficient Parts-Based NearDuplicate and Sub-Image Retrieval System," *ACM International Conference on Multimedia*, pp **869–876**.

45. LiG., WuQ., TuD., and SunS.,(2007) "A Sorted Neighborhood Approach for Detecting Duplicated Regions in Image Forgeries Based on DWT and SVD," IEEE International Conference on Multimedia and Expo, pp **1750–1753**.
46. LangilleA. and GongM., (2006) "An Efficient Match-based Duplication Detection Algorithm," Canadian Conference on Computer and Robot Vision, pp **64–71**.
47. LinC.Y., WuM., BloomJ., CoxI., MillerM., and LuiY., (2001) "Rotation, Scale, and Translation Resilient Watermarking for Images," IEEE Transactions on Image Processing, 10, 5, pp **767–782**.
48. LinH., WangC., and KaoY., (2009) "Fast Copy-Move Forgery Detection," WSEAS Transactions on Signal Processing, 5, 5, pp **188–197**.
49. LuoW., HuangJ., and QiuG., (2006) "Robust Detection of Region-Duplication Forgery in Digital Images," in International Conference on Pattern Recognition, 4, pp **746–749**.
50. MahdianB. and SaicS., (2007) "Detection of Copy-Move Forgery using a Method Based on Blur Moment Invariants," Forensic Science International, 171, 2, pp **180–189**.
51. MyrnaA. N., VenkateshmurthyM. G., and PatilC. G., (2007) "Detection of Region Duplication Forgery in Digital Images Using Wavelets and LogPolar Mapping," IEEE International Conference on Computational Intelligence and Multimedia Applications, pp **371–377**.
52. PopescuA. and FaridH., (2004)"Exposing digital forgeries by detecting duplicated image regions," Department of Computer Science, Dartmouth College, Tech. Rep. TR2004-515. Available: www.cs.dartmouth.edu/farid/publications/tr04.html
53. RyuS., LeeM., and LeeH., (2010) "Detection of Copy-Rotate-Move Forgery using Zernike Moments," in Information Hiding Conference, pp **51–65**.
54. ShiehJ.M., LouD.C., and ChangM.C., (2006) "A Semi-Blind Digital Watermarking Scheme based on Singular Value Decomposition," Computer Standards & Interfaces, 28, 4, pp **428–440**.
55. WangJ., LiuG., ZhangZ., DaiY., and WangZ., (2009) "Fast and Robust Forensics for Image Region-Duplication Forgery," ActaAutomaticaSinica, 35, 12, pp **1488–1495**.
56. WangJ., LiuG., LiH., DaiY., and WangZ., (2009) "Detection of Image Region Duplication Forgery Using Model with Circle Block" International Conference on Multimedia Information Networking and Security, pp **25–29**.
57. ZhangJ., FengZ., and SuY., (2008) "A New Approach for Detecting CopyMove Forgery in Digital Images," in International Conference on Communication Systems, pp **362–366**.

58. AkbarpourSekeh M., AizainiMaarof M., FoadRohani M., Mahdian B., (2013) "Efficient image duplicated region detection model using sequential block clustering", Digital Investigation, 10, pp 73–84.
59. PopescuA. and FaridH., (2004) "Exposing digital forgeries by detecting duplicated image regions," Department of Computer Science, Dartmouth College, Tech. Rep. TR2004-515. Available: www.cs.dartmouth.edu/farid/publications/tr04.html
60. HuangH., GuoW., and ZhangY., (2008) "Detection of Copy-Move Forgery in Digital Images Using SIFT Algorithm," Pacific-Asia Workshop on Computational Intelligence and Industrial Application, 2, pp 272–276.
61. ChristleinV., RiessC., and AngelopoulouE., (2010) "A Study on Features for the Detection of Copy-Move Forgeries," in GI SICHERHEIT.
62. PanX. and LyuS., (2010) "Region Duplication Detection Using Image Feature Matching," IEEE Transactions on Information Forensics and Security, 5, 4, pp 857–867.
63. MahdianB. and SaicS., (2007) "Detection of Copy-Move Forgery using a Method Based on Blur Moment Invariants," Forensic Science International, 171, 2, pp 180–189.
64. Bayram S., Sencar H.T., Memon N., (2010) "A SURVEY OF COPY-MOVE FORGERY DETECTION TECHNIQUES," IEEE Western New York Image Processing Workshop.
65. Cao Y., Gao T., Fan L. , Yang Q., (2012) "A robust detection algorithm for copy-move forgery in digital images"Forensic Science International, 214, 3, pp 33–43.
66. TurkM. and PentlandA., (1991) "Eigenfaces for recognition", Journal of Cognitive Neuroscience, 3, 1, pp 201-209.
67. Shao H., Yu T., Xu M., Cui W., (2012) "Image region duplication detection based on circular window expansion and phase correlation", 222, pp 71–82.
68. Bay H., Ess A., Van Gool L., (2008) "Speeded-Up Robust Features (SURF)" , Computer Vision and Image Understanding 110, pp 346–359.
69. ShuaiX., ZhangC., and HaoP., (2008) "Fingerprint indexing based on composite set of reduced SIFT features," in Proc. of ICPR, Tampa, Florida, USA.
70. AmeriniI., BallanL., CaldelliR., Del BimboA., and SerraG., (2010) "Geometric tampering estimation by means of a SIFT-based forensic analysis," in Proc. of IEEE ICASSP, Dallas, USA.

71. Gajanan K., Vijay H., (2013) “Digital image forgery detection using passive techniques”, Digital Investigation, pp **1–20**.
72. Lowe D.G., (2004) “Distinctive Image Features from Scale-Invariant Keypoints”, IJCV, 60, pp **91-110**.
73. Hu M.K., (1962) “Visual pattern recognition by moment invariants” IEEE Trans. Information Theory 8, pp **179–187**.
74. Kim H.S., and Lee H.K., (2003) “Invariant image watermark using Zernike moments. IEEE Trans. Circuits and Systems for Video Technology 13, 8, pp **766–775**.
75. Teh C.H., and Chin R.T., (1988) “On image analysis by the methods of moments”. IEEE, Trans. Pattern Analysis and Machine Intelligence 10, 4, pp **496–513**.
76. Khotanzad A., and Hong Y.H., (1990), “Invariant image recognition by Zernike moments”, IEEE Trans. Pattern Analysis and Machine Intelligence 12, 5, pp **489–497**.
77. Zernike F., (1934) “Beugungstheorie des schneidenverfahrens und seiner verbessertenformderphasenkontrastmethode”. Physica 1, pp **689–704**.
78. HastieT., TibshiraniR., and FriedmanJ. H., (2003) “The Elements of Statistical Learning”. Springer.
79. HartleyR. I. and ZissermanA., (2004) “Multiple View Geometry in Computer Vision”. Cambridge University Press.
80. FischlerM. and BollesR., (1981) “Random sample consensus: a paradigm for model fitting with applications to image analysis and automated cartography,” Communications of the ACM, 24, 6, pp **207-215**.
81. NgT.T., ChangS.F., HsuJ., and PepeljugoskiM., (2004) “Columbia photographic images and photorealistic computer graphics dataset,” ADVENT,Columbia University, Tech. Rep.
82. Huang Y., Lu W., Sun W., and Long D., (2011) “improved dct-based detection of copy-move forgery in images”, Forensic Science International, 206, pp **178-184**.

Abstract

In this thesis we propose an efficient and fully automatic method to detect copy -move forgery detection in digital images.

Nowadays wide use of the Internet and digital images in it, have made such images sources of information, in courts of law they may be used as testimony to proof or refuse evidences, and also a number of tampered images have been published by newspapers. So determining image originality is a fundamental task.

There are several methods to make fake images, but the most common is copy-move forgery, that the forger copies a part(s) of image and pasted it into another part(s) of that image. Many researchers have done beneficial researches on it. Most of them can find only copy-moved forgery. In other words, they can find regions which are only copied and pasted without any changes and are weak against some types of manipulations. Many forgers make some changes on copied regions so that the image sounds more natural.

So a good copy-move forgery detection technique should be robust to some types of changes such as rotation, scaling, JPEG compression, Gaussian noise addition and gamma correction. Most existing methods do not deal with all these manipulations and often have time complexity.

Detection method based on SIFT features is robust to all such modifications, but failed to find flat copied regions.

Zernike moments are invariant against rotation, additive white Gaussian noise, JPEG compression, and blurring. They can find flat copied regions too, but sensitive to scaling. So it is clear that using these two features is very proper to detect all copied regions in an image.

So our scheme is robust against scaling, rotation, additive white Gaussian noise, JPEG compression, gamma correction and affine transformation and the following improvements have been imposed on it:

- ✓ Detecting more than one image duplicated region from a single region, while other methods can find only a single copy of a region.
- ✓ It can distinguish between similar object and regions, from copied objects and regions.

So by combining detection methods based on SIFT features and Zernike moments, the proposed method is robust against flat copied region in addition to mentioned modifications. Further more In the case of fraud detection, geometric transformation parameters are determined. Our proposed method is precise and has a reduced computational complexity and time processing.



Shahrood University of Technology
Faculty Computer Engineering & Information Technology
M. Sc. Thesis

Proposing an efficient method to detect forged digital images

Zahra Mohamadian

Supervisor:

Dr. Ali Akbar Pouyan

Advisor:

Dr. Hamid Hassan pour

August 2013