

نهان نگاری تصویر: مروری بر پیشرفت های اخیر^۱

نویسندگان:

NANDHINI SUBRAMANIAN² (Member, IEEE)
OMAR ELHARROUSS²,
SOMAYA AL-MAADEED² (Senior Member, IEEE),
AHMED BOURIDANE³ (Senior Member, IEEE)

¹ Image Steganography: A Review of the Recent Advances

این مقاله در ۲۵ ژانویه ۲۰۲۱ در IEEE Access منتشر شده است:

Subramanian, N., Elharrouss, O., Al-Maadeed, S., & Bouridane, A. (2021). Image steganography: A review of the recent advances. *IEEE Access*.

² Department of Computer Science and Engineering, Qatar University, Doha, Qatar

³ Department of Computer and Information Sciences, Northumbria University, Newcastle upon Tyne NE1 8ST, U.K.

چکیده

نهان‌نگاری تصویر، فرآیند پنهان کردن اطلاعاتی است که می‌تواند متن، تصویر یا ویدیو در داخل یک تصویر رویه^۱ باشد. اطلاعات محرمانه به گونه‌ای پنهان است که برای چشم انسان قابل مشاهده نیست. فناوری یادگیری عمیق که به عنوان یک ابزار قدرتمند در کاربردهای مختلف از جمله نهان‌نگاری تصویر ظاهر شده است، اخیراً مورد توجه قرار گرفته است. هدف اصلی این مقاله بررسی و بحث در مورد روش‌های یادگیری عمیق موجود در زمینه نهان‌نگاری تصویر است. تکنیک‌های یادگیری عمیق مورد استفاده برای نهان‌نگاری تصویر را می‌توان به طور کلی به سه دسته تقسیم کرد: روش‌های سنتی، روش‌های مبتنی بر شبکه عصبی کانولوشنی^۲ و روش‌های مبتنی بر شبکه مولد تخصصی^۳. همراه با روش، خلاصه‌ای مفصل در مورد مجموعه داده‌های مورد استفاده، مجموعه‌های آزمایشی در نظر گرفته شده و معیارهای ارزیابی که معمولاً استفاده می‌شود، در این مقاله توضیح داده شده است. جدولی که تمام جزئیات را خلاصه کند نیز برای مراجعه آسان، ارائه شده است. هدف این مقاله، گردآوری روندها، چالش‌ها و برخی جهت‌گیری‌های آینده در این زمینه و کمک به محققین همکار است.

کلمات کلیدی:

Image steganography, GAN steganography, CNN steganography, information hiding, image data hiding.

¹ Cover Image

² Convolutional Neural Networks

³ Generative Adversarial Network-based

۱- مقدمه

فناوری در طول سال‌های گذشته گسترش یافته و منجر به استفاده گسترده از چند رسانه‌ای برای انتقال داده‌ها، به ویژه اینترنت اشیا (IoT) شده است. معمولاً انتقال از طریق کانال‌های شبکه ناامن انجام می‌شود. به طور خاص، اینترنت برای تبادل رسانه‌های دیجیتال محبوبیت شتابانی به دست آورده است و افراد، شرکت‌های خصوصی، موسسات، دولت‌ها از این روش‌های انتقال داده‌های چندرسانه‌ای برای تبادل داده‌ها بهره می‌برند. اگرچه مزایای متعددی به آن پیوست شده است، اما یکی از معایب برجسته آن، حفظ حریم خصوصی و امنیت داده‌ها است.

در دسترس بودن ابزارهای متعددی که به راحتی در دسترس هستند و قادر به بهره‌برداری از حریم خصوصی، یکپارچگی داده‌ها و امنیت داده‌های در حال انتقال هستند، امکان تهدیدهای مخرب، استراق سمع و سایر فعالیت‌های خرابکارانه را ایجاد کرده است.

راه حل برجسته، رمزگذاری داده است که در آن، داده‌ها با استفاده از کلید رمزگذاری به یک دامنه متن رمزی تبدیل می‌شوند. در انتهای دریافت، متن رمز شده با استفاده از یک کلید رمزگشایی به متن ساده تبدیل می‌شود. در استفاده از رمزگذاری داده‌ها، داده‌های اصلی قابل مشاهده نیستند، با این حال، متن رمز شده به شکل درهم به چشم انسان قابل مشاهده است که منجر به سوءظن و بررسی بیشتر می‌شود. یک موضوع تحقیقاتی جدید، نهان‌نگاری، در این زمینه مورد قبول قرار گرفته تا داده‌هایی را که برای چشم انسان قابل درک نیست را پنهان کند.

تکنیک‌های پنهان کردن اطلاعات برای مدت طولانی در دسترس بوده‌اند اما اهمیت آنها اخیراً افزایش یافته است. دلیل اصلی، افزایش ترافیک داده‌ها از طریق اینترنت و شبکه‌های اجتماعی است. اگرچه اهداف رمزنگاری و نهان‌نگاری مشابه هستند، اما تفاوت ظریفی وجود دارد. رمزنگاری، داده‌ها را ناشکستنی و غیرقابل خواندن می‌کند؛ اما نهان‌نگاری برای چشم انسان قابل مشاهده است.

نهان‌نگاری، که برای پنهان کردن اطلاعات در دید ساده استفاده می‌شود، امکان استفاده از طیف گسترده‌ای از اشکال اطلاعات مخفی مانند تصویر، متن، صدا، ویدئو و فایل‌ها را فراهم می‌کند. واترمارک دیجیتال^۱ روش دیگری است که در آن اطلاعات محرمانه برای ادعای مالکیت جاسازی می‌شود. رمزنگاری روش محبوبی است که برای پنهان کردن اطلاعات استفاده می‌شود، اما، نهان‌نگاری در دوران اخیر محبوبیت پیدا کرده است.

نهان‌نگاری را می‌توان به عنوان فرآیند پنهان کردن یک داده چند رسانه‌ای کوچک در داخل داده‌های چند رسانه‌ای دیگر اما بسیار بزرگتر مانند تصویر، متن، فایل یا ویدئو تعریف کرد [۱]. نهان‌نگاری تصویر، تکنیکی

^۱ Digital Watermark

برای پنهان کردن یک تصویر در داخل یک تصویر دیگر است. در نهان‌نگاری تصویر، تصویر رویه به گونه‌ای دستکاری می‌شود که داده‌های پنهان، دیده نمی‌شوند و در نتیجه مانند مورد رمزنگاری، مشکوک نمی‌شوند.

برعکس، نهان‌واری^۱ برای تشخیص وجود هر پیام مخفی پوشیده شده در تصویر و استخراج داده‌های پنهان استفاده می‌شود [۲]. نهان‌واری به طبقه‌بندی اینکه آیا تصویر، یک تصویر نهانی یا یک تصویر معمولی است، کمک می‌کند. جدای از طبقه‌بندی تصویر، تحقیقات بیشتری برای شناسایی مکان و محتوای تصویر مخفی در داخل تصویر رویه انجام می‌شود.

با در دسترس بودن حجم عظیمی از داده‌ها، یادگیری عمیق (DL) به یک روند تبدیل شده است و به طور گسترده برای بسیاری از برنامه‌ها استفاده می‌شود. یادگیری عمیق یک ابزار مفید در کاربردهای مختلف مانند طبقه‌بندی تصویر، تشخیص خودکار گفتار، تشخیص تصویر، پردازش زبان طبیعی، سیستم‌های توصیه، پردازش تصاویر پزشکی است [۳].

اگرچه تحقیقات در مورد نهان‌نگاری کاملاً جدید است، اما از روش‌های DL از جمله روش‌های مبتنی بر شبکه‌های عصبی کانولوشنی (CNN) شبکه‌های مولد تخاصمی (GANs) و استقرار آنها در هر دو وجه نهان‌نگاری و نهان‌واری استفاده شده است.

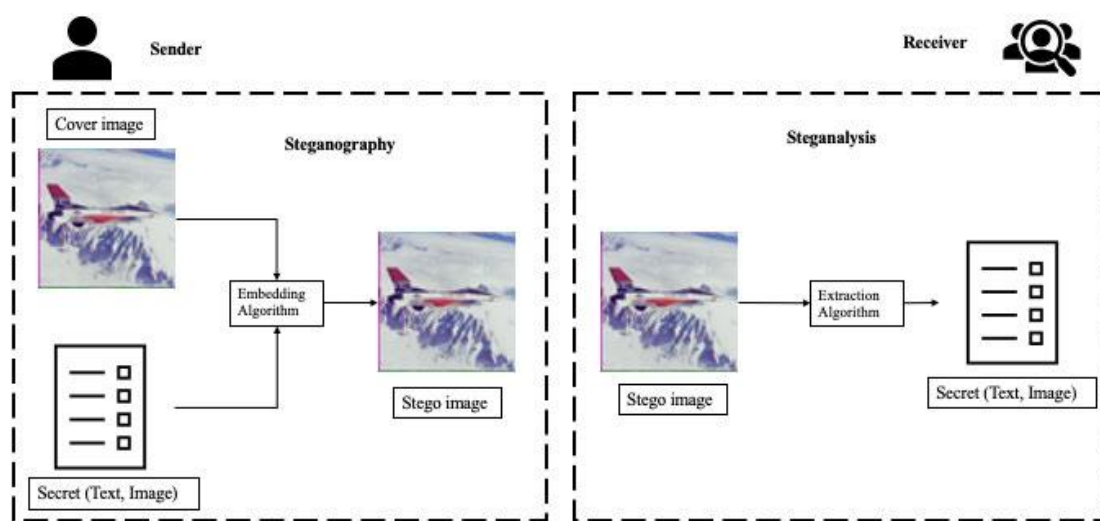
هدف اصلی این مقاله، بررسی روش‌های موجود، ارائه روندها و بحث در مورد چالش‌هایی است که در حال حاضر در مطالعات موجود است. همراه با این مطالعات، مجموعه داده‌هایی که در دسترس عموم هستند و معمولاً مورد استفاده قرار می‌گیرند و معیارهای ارزیابی در نظر گرفته شده نیز مورد بحث قرار می‌گیرند. در نهایت، مقایسه عملکرد بین روش‌ها و بحث احتمالی برای شناسایی شکاف‌ها در مطالعات حاضر، مزایا و معایب روش‌ها تشریح شده است.

ادامه مقاله به شرح ذیل سازماندهی شده است. بخش دوم اصل کار روش‌ها را به سه دسته خلاصه می‌کند (روش‌های سنتی، روش‌های مبتنی بر CNN و روش‌های مبتنی بر GAN). مجموعه داده‌هایی که معمولاً استفاده می‌شوند در بخش ۳ همراه با ارزیابی در بخش ۴ توضیح داده شده‌اند. جدولی با مقایسه نتایج حاصل از روش‌های مختلف با مجموعه‌های آزمایشی که عموماً در بخش ۴ استفاده می‌شوند ارائه شده است. در نهایت، چالش‌های پیش‌رو، یک بحث مختصر و نتیجه‌گیری به ترتیب در بخش ۵، ۶ و بخش ۷ اضافه شده است.

^۱ Steganalysis

۲- مختصری از روش‌ها

پس از بررسی تمام چارچوب‌های موجود، روش‌ها در درجه اول به سه دسته تقسیم می‌شوند، یعنی روش‌های سنتی نهان‌نگاری تصویر، روش‌های نهان‌نگاری تصویر مبتنی بر CNN و روش‌های نهان‌نگاری تصویر مبتنی بر GAN. روش‌های سنتی چارچوب‌هایی هستند که از روش‌هایی استفاده می‌کنند که به یادگیری ماشین یا الگوریتم‌های یادگیری عمیق مرتبط نیستند. بسیاری از روش‌های سنتی مبتنی بر تکنیک LSB هستند. روش‌های مبتنی بر CNN مبتنی بر شبکه‌های عصبی کانولوشنی عمیق برای جاسازی و استخراج پیام‌های مخفی هستند و روش‌های مبتنی بر GAN از برخی از انواع GAN استفاده می‌کنند. شکل ۱، یک نمای کلی از معماری نهان‌نگاری و نهان‌واری را نشان می‌دهد. همانطور که در شکل ۱ نشان داده شده است، ورودی‌ها، تصویر رویه و اطلاعات مخفی هستند که می‌توانند متن یا تصویر باشند. مدل DL می‌تواند مبتنی بر CNN یا مبتنی بر GAN باشد. در حالی که بلوک نهان‌نگاری تصویر پوشش^۱ را تولید می‌کند، مدل نهان‌واری تصویر پوشش را به عنوان ورودی برای شناسایی و شاید استخراج اطلاعات مخفی، می‌گیرد. در برخی از روش‌ها، احتمال عادی بودن تصویر ورودی یا تصویر پوشش به عنوان خروجی، محاسبه می‌شود.



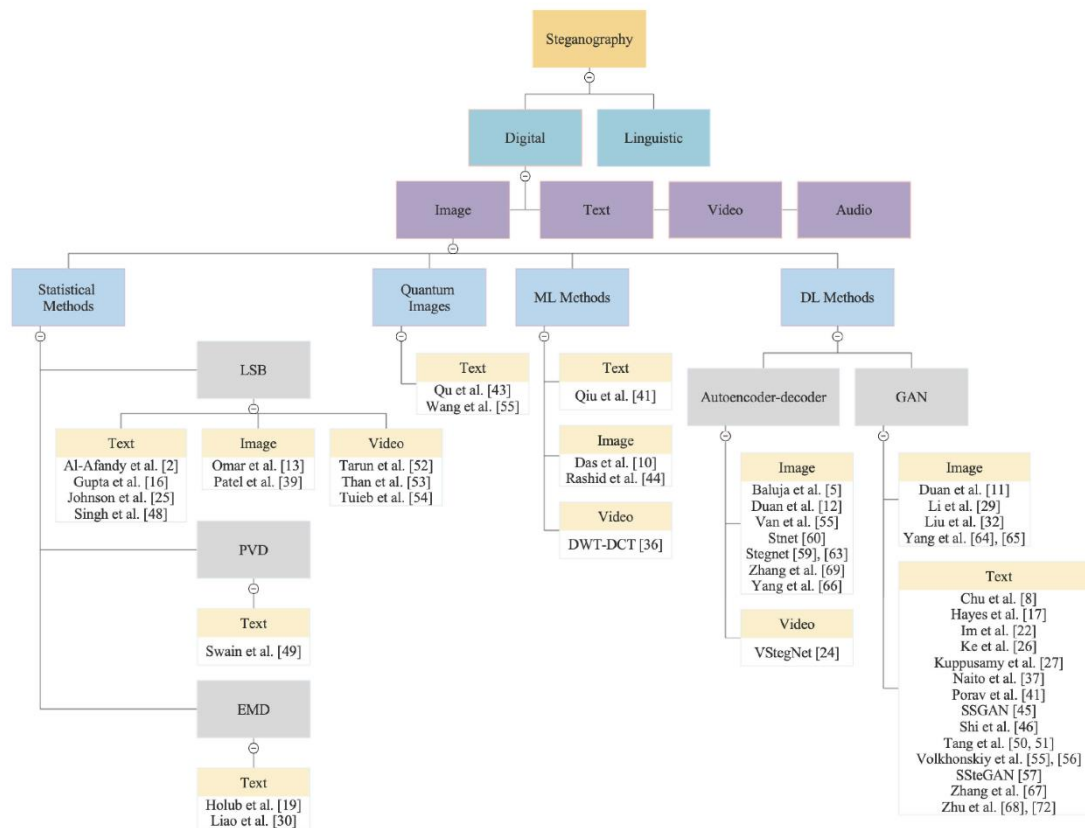
شکل ۱- اصل کار کلی نهان‌نگاری و نهان‌واری. ورودی‌ها، تصویر رویه و اطلاعات مخفی هستند و از یک الگوریتم جاسازی برای تولید حامل تصویر نهان‌نگاری شده استفاده می‌شود. الگوریتم استخراج، تصویر نهان‌نگاری شده را به عنوان ورودی برای استخراج اطلاعات محرمانه، می‌گیرد.

داده‌های متنی، تصاویر رنگی یا خاکستری معمولاً به عنوان رسانه مخفی استفاده می‌شوند. دو عامل - ماهیت رسانه مخفی و تکنیک مورد استفاده - در طبقه بندی آثار موجود در دسته‌های مختلف در نظر گرفته شده است که در شکل ۲ مشاهده می‌شود. از تجزیه و تحلیل در شکل ۲، مشاهده می‌شود که متن^۲،

^۱ Stego image

^۲ Text

رایجترین اطلاعات محرمانه مورد استفاده است و روش‌های مبتنی بر GAN حالت ترجیحی برای ارتباط مخفی برای پنهان کردن متن هستند.



شکل ۲- طبقه‌بندی روش‌های موجود بر اساس رسانه‌های مخفی و روش مورد استفاده. ابتدا روش‌ها طبقه‌بندی شده و در زیر هر دسته، طبقه‌بندی بیشتر بر اساس رسانه‌های مخفی انجام شده است.

۱-۲ نمان‌نگاری مبتنی بر روش‌های سنتی^۱

به طور مرسوم، روش جایگزینی حداقل بیت‌های مهم (LSB) برای انجام نمان‌نگاری تصویر استفاده می‌شود. تصاویر معمولاً دارای کیفیت پیکسل بالاتری هستند که از همه پیکسل‌ها استفاده نمی‌شود. روش‌های LSB با این فرض کار می‌کنند که اصلاح چند مقدار پیکسل، هیچ تغییر قابل مشاهده‌ای را نشان نمی‌دهد. اطلاعات مخفی به یک فرم باینری یا دودویی تبدیل می‌شود. تصویر پوششی پویش^۲ می‌شود تا LSB‌ها در ناحیه نویزی را مشخص کند. سپس بیت‌های باینری از تصویر مخفی در LSB‌های تصویر پوشش جایگزین می‌شوند. روش جایگزینی باید با احتیاط انجام شود زیرا سربارگذاری^۳ تصویر پوشش ممکن است منجر به تغییرات قابل مشاهده شود که آن، منتهی به نشت حضور اطلاعات مخفی می‌شود [۴] و [۵].

^۱ Traditional-Based

^۲ Scanning

^۳ overloading

با روش LSB به عنوان خط پایه^۱، تعدادی از روش‌های مرتبط پیشنهاد شده است. به عنوان مثال، یک تغییر جزئی در تبدیل پیام مخفی به کدهای دودویی در [۶] پیشنهاد شده است. ابتدا یک روش رمزگذاری هافمن^۲ برای رمزگذاری پیام مخفی در بیت‌های دودویی به کار گرفته می‌شود. سپس، بیت‌های کدگذاری شده با استفاده از روش LSB در تصویر پوشش جاسازی می‌شوند. در [۷] نسخه دیگری از روش LSB برای تصاویر RGB استفاده شده است. تصویر پوشش ۳ کانالی است و برش بیتی داده شده است. پیام مخفی در هر سه صفحه به نسبت ۲:۲:۴ برای سطوح R، G و B تعبیه شده است. نه تنها از دامنه مکانی^۳، بلکه از تصاویر کوانتومی^۴ نیز استفاده می‌شود [۸] و [۹]. دامنه فرکانس^۵ در دامنه تصویر کوانتومی مورد استفاده قرار می‌گیرد و پیکسل‌هایی که به نظر می‌رسد بر رنگ تأثیر می‌گذارند برای پنهان کردن بیت‌های مخفی استفاده می‌شوند. ترکیبی از رمزنگاری و نهان‌نگاری در جایی استفاده می‌شود که LSB تصویر پوشش با مهم‌ترین بیت‌های تصویر مخفی جایگزین می‌شود [۱۰]. مولد اعداد شبه تصادفی برای انتخاب پیکسل‌ها مورد استفاده قرار می‌گیرد و هر بار با استفاده از چرخش، کلید رمزگذاری می‌شود. یک روش k-LSB پیشنهاد شده است که در آن k عدد از کمترین بیت با پیام مخفی [۱۱] جایگزین می‌شود. برای نهان‌واری، از یک فیلتر آنتروپی^۶ برای شناسایی و استخراج تصویر مخفی استفاده می‌شود [۱۱].

روش‌های LSB برای مخفی کردن اطلاعات محرمانه در داخل فیلم‌ها نیز استفاده می‌شود. ویدئوها دنباله‌ای از تصاویر هستند که به آن فریم‌های ویدئویی^۷ می‌گویند. هر ویدئو در فریم‌های تصویر تجزیه می‌شود و بیت‌های باینری اطلاعات مخفی در LSB فریم‌های تصویر ویدئو پنهان می‌شود. یک شکل اساسی از روش جایگزینی LSB [۱۲] و ترکیبی از رمزگذاری هافمن و روش‌های جایگزینی LSB در ویدئوها استفاده می‌شود [۱۳]. یکی دیگر از رویکردهای جالب این است که در کنار فریم‌های تصویر ویدئو، از صدا نیز برای تقویت پنهان‌سازی استفاده می‌شود [۱۴]. علاوه بر روش‌های LSB [۱۵]، ترکیبی از تبدیل کسینوس گسسته^۸ (DCT) و تبدیل موجک گسسته^۹ (DWT) را برای پنهان کردن پیام مخفی در داخل یک ویدیوی پوشش^{۱۰} پیشنهاد کرده است. برای یافتن مناطق مورد نظر، از روش ردیابی شیء چندگانه^{۱۱} (MOT) استفاده می‌شود. داده‌های مخفی ابتدا کدگذاری می‌شوند و سپس به بیت‌های باینری تبدیل می‌شوند و در نهایت، آنها را در ویدیوی پوشش، جاسازی می‌کنند.

جدول ۶ مروری دقیق بر روش‌های سنتی نهان‌نگاری تصویر را نشان می‌دهد. یکی دیگر از روش‌های کلاسیک مورد استفاده در زمینه نهان‌نگاری تصویر، تفاوت ارزش پیکسل^{۱۲} (PVD) است. PVD با گرفتن اختلاف

¹ Baseline

² Huffman encoding method

³ Spatial

⁴ Quantum images

⁵ Frequency domain

⁶ Entropy filter

⁷ video frames

⁸ Discrete Cosine Transformation

⁹ Discrete Wavelet Transformation

¹⁰ cover video

¹¹ multiple objects tracking method

¹² Pixel Value Differencing

پیکسل‌های متوالی برای یافتن مکان‌های پنهان کردن بیت‌های مخفی کار می‌کند، به گونه‌ای که یکنواختی تصویر رویه حفظ شود. برای هر ۸ بیت، ترکیبی از LSB در دو بیت اول و PVD در شش بیت باقی مانده طراحی شده است [۱۸]. علاوه بر این، برخی از تکنیک‌های دیگر مورد استفاده، نهان‌نگاری بدون پوشش است که در آن تصویر پوشش ارائه نمی‌شود، بلکه بر اساس اطلاعات مخفی تولید می‌شود. اطلاعات مخفی گرفته می‌شود و یک مدیریت ارتباط برای تولید تصویر پوشش با استفاده از روش تشخیص شیء انجام می‌شود. نهان‌نگاری بدون پوشش مشابهی پیشنهاد شده که در آن، ویژگی‌های الگوهای باینری محلی^۱ (LBP) تصویر پوشش و تصاویر مخفی ابتدا هش می‌شوند. بعداً، هش‌ها برای ایجاد تصویر نهان‌نگاری شده، مطابقت داده می‌شوند [۱۹]. به همین ترتیب، به جای LBP، لبه‌های تصاویر پوشش رنگی به دست می‌آید. سپس، بیت‌های باینری اطلاعات مخفی در لبه‌های کشف شده در تصاویر پوشش، پنهان می‌شوند [۲۰].

جدول اخلاص‌های از جزئیات روش‌های سنتی.

روش	مجموعه داده‌ها	معیارها	مزایا	معایب
[۱۶]	1 RGB Image	PSNR and Time	- زمان کم محاسبه - قوی در وارد کردن و استخراج - نهان‌نگاری و نهان‌واری بدون وابستگی می‌توانند کار کنند	- امنیت کمتر - اطلاعات مخفی صرفاً متنی است
[۱۷]	Lena and Baboon	PSNR and MSE	- زمان کم محاسبه - پیام مخفی، تصویر است - فورمت دلخواه تصاویر مورد قبول است	- امنیت در مقایسه با روش‌های یادگیری عمیق، کمتر است
[۱۰]	Lena	PSNR	- زمان کم محاسبه - پیام مخفی، تصویر است	- امنیت کم

انواع دیگر روش‌های سنتی نهان‌نگاری در ادامه شرح داده شده است. از تصاویر پزشکی JPEG بهره‌گیری شده و عمل جاسازی با در نظر گرفتن تفاوت در ضرایب DCT پیکسل‌ها بین دو بلوک متوالی با در نظر گرفتن پیکسل‌ها در همان موقعیت‌ها در دو بلوک انجام می‌شود [۲۱]. یک روش جدید به نام هیستوگرام چگالی پیکسل (PHD) برای تصاویر نیمه تمام پیشنهاد شده است [۲۲]. چگالی پیکسل تصاویر محاسبه می‌شود و یک هیستوگرام چگالی پیکسل برای پنهان کردن اطلاعات مخفی تشکیل می‌شود. از توزیع پواسون برای به دست آوردن خطای پشت سر هم استفاده می‌شود و بازسازی تصاویری که مقاوم در برابر فشرده سازی هستند با استفاده از رمزگشایی STC انجام می‌شود [۲۳]. برای توضیحات بیشتر در مورد روش‌های سنتی نهان‌نگاری، [۲۴] قابل ارجاع است.

^۱ Local Binary Patterns

۲-۲- روش‌های نهان‌نگاری مبتنی بر شبکه‌های عصبی کانولوشنی (CNN)^۱

در نهان‌نگاری تصویر با استفاده از مدل‌های CNN به شدت از معماری رمزگذار-رمزگشا^۲ الهام گرفته شده است. دو ورودی (تصویر رویه و تصویر مخفی) به عنوان ورودی به رمزگذار برای تولید تصویر پوشش استفاده شده و تصویر پوشش به عنوان ورودی به رمزگشا برای استخراج تصویر مخفی جاسازی شده، داده می‌شود. اصل اساسی یکسان است با این تفاوت که روش‌های مختلف، معماری‌های مختلفی را امتحان کرده‌اند. نحوه الحاق تصویر رویه و ورودی و تصویر مخفی نیز در رویکردهای مختلف، متفاوت است؛ در حالی که تغییرات در لایه کانولوشن، لایه ادغام مورد انتظار است. تعداد فیلترهای استفاده شده، گام‌ها، اندازه فیلتر، عملکرد فعال‌سازی استفاده شده و عملکرد از دست دادن از روشی به روش دیگر متفاوت است. یک نکته مهم که در اینجا باید به آن توجه کرد، اندازه تصویر رویه و تصویر مخفی باید یکسان باشد، بنابراین هر پیکسل از تصویر مخفی در تصویر رویه توزیع می‌شود.

وو و همکاران^۳ [۲۵] و [۲۶] یک معماری رمزگذار-رمزگشا را پیشنهاد کرده‌اند. معماری رمزگذار-رمزگشا مبتنی بر U-Net برای مخفی کردن استفاده می‌شود و یک CNN با ۶ لایه برای استخراج، در [۲۷] توسط دوان^۴ و همکاران پیشنهاد شده است. شکل ورودی U-Net برای پذیرش 256*256 و ۶ کانال اصلاح شده است. تصاویر مخفی و رویه تا ورودی به هم پیوسته‌اند و از این روی، ۶ کانال ارائه شود. یک شبکه پنهان (H-net) و آشکارساز (R-net) مبتنی بر U-net توسط ون و همکاران^۵ استفاده می‌شود. در [۲۸]. نرمال‌سازی دسته‌ای^۶ و فعال‌ساز ReLU استفاده می‌شود. تصاویر رویه و مخفی قبل از ارسال به شبکه به هم متصل می‌شوند. دو بهینه‌ساز خطای SSIM و MSE برای کاهش خطاها و در نتیجه بهبود عملکرد، استفاده می‌شود.

یک کانولوشن جداپذیر با بلوک باقیمانده (SCR) برای به هم پیوستن تصویر رویه و تصویر مخفی استفاده می‌شود [۲۵]. تصویر تعبیه شده به عنوان ورودی به رمزگذار برای ساخت تصویر پوشش داده می‌شود که برای استخراج تصویر رمزگشایی شده به رمزگشا داده می‌شود. ELU (واحد خطی نمایی) و نرمال‌سازی دسته‌ای استفاده می‌شود. یک تابع هزینه جدید برای کاهش اثر نویز در تصویر حامل تولید شده به نام تلفات واریانس پیشنهاد شده است [۲۶].

یک معماری رمزگذار-رمزگشا توسط رحیم و همکاران^۷ ارائه شده است. در [۲۹]، این روش با سایر روش‌ها در نحوه ارائه ورودی‌ها متفاوت است. بخش رمزگذار شامل دو معماری موازی برای پوشش و تصویر مخفی است. ویژگی‌های تصویر رویه و تصاویر مخفی از طریق لایه کانولوشن استخراج شده و به هم متصل می‌شوند. ویژگی‌های به هم پیوسته برای ساختن تصویر پوشش مورد بهره‌برداری قرار می‌گیرد.

¹ CNN-BASED STEGANOGRAPHY METHODS

² encoder-decoder

³ Wu et al.

⁴ Duan

⁵ Van et al.

⁶ Batch Normalization

⁷ Rahim et al.

رویکردی کمی متفاوت، توسط وانگ و همکاران^۱ در [۳۰] با در نظر گرفتن حالت ظاهری عکس همراه با اطلاعات مخفی و تصویر رویه پیشنهاد شده است. تصویر پوششی ایجاد شده به تصویر سبک^۲ داده شده به عنوان ورودی، تبدیل می‌شود. شبکه آشکار برای رمزگشایی اطلاعات مخفی از تصویر پوششی ایجاد شده، استفاده می‌کند. مشابه روش‌های دیگر، معماری رمزگذار-رمزگشا خودکار با VGG به عنوان پایه استفاده شده است. اندازه تصویر دلخواه برای اطلاعات مخفی و سبک تصویر با استفاده از لایه عادی سازی نمونه تطبیقی (AdaIN) گرفته می‌شود و خروجی اندازه تصویر رویه است. توزیع پیکسلی تصویر رویه با استفاده از یک پیکسل CNN توسط یانگ و همکاران در [۳۱] بدست می‌آید. سپس، اطلاعات مخفی با کاهش نمونه‌برداری در توزیع پیکسل، به طور یکنواخت جاسازی می‌شود.

سه شبکه یعنی شبکه آماده‌سازی^۳، شبکه پنهان‌سازی^۴ و شبکه آشکارسازی^۵ توسط بالوجا و همکاران^۶ پیشنهاد شده است. در [۳۲] بر اساس معماری رمزگذار خودکار به عنوان یک مدل واحد. از شبکه آماده‌سازی برای تهیه تصویر مخفی قبل از تغذیه آن به عنوان ورودی به شبکه پنهان استفاده می‌شود، که خروجی شبکه آماده‌سازی و تصویر پوششی را برای تولید تصویر، می‌گیرد. شبکه آشکارسازی تصویر مخفی را از تصویر حامل، با کشف تصویر رویه، رمزگشایی می‌کند. دو مقدار تلفات بین رویه و حامل ساخته شده و بین تصویر مخفی و رمزگشایی شده محاسبه می‌شود. مدل با استفاده از شاخص تشابه ساختار (SSIM) ارزیابی شده است. توسعه‌ای در [۳۲] توسط ژانگ و همکاران پیشنهاد شده است. در [۳۳] که در آن تصویر رویه به فرمت تصویر $YCrCb$ تبدیل می‌شود و تصویر مخفی فقط در کانال Y پنهان می‌شود زیرا تمام اطلاعات معنایی و رنگی در کانال‌های C_r و C_b وجود دارد. همچنین، برای کاهش دو سوم بار، تصویر مخفی به فرمت تصویر در مقیاس خاکستری تبدیل می‌شود. کانال Y تصویر رویه و تصویر مخفی در مقیاس خاکستری به عنوان ورودی به شبکه رمزگذار-رمزگشا برای ساخت تصویر پوششی داده می‌شود. تصویر پوششی خروجی کانال Y است که با کانال‌های C_r و C_b تصویر رویه ترکیب می‌شود تا تصویر پوششی را در فضای رنگی $YCrCb$ ایجاد کند. برای استخراج تصویر مخفی، کانال Y تصویر پوششی، دوباره به شبکه آشکار ساز داده می‌شود تا تصویر مخفی در مقیاس خاکستری را خروجی کند. دو نوع مختلف (مدل‌های پایه و باقیمانده) نیز در مدل‌های مولد استفاده می‌شود. یک تابع خطای مختلط که برای نهان‌نگاری با استفاده از SSIM مناسب‌تر است و شاخص تشابه ساختار چند مقیاسی آن (MS-SSIM) نیز استفاده می‌شود. برای شبکه پنهان سازی از معماری ISGAN استفاده شده است. جدول ۲ بررسی روش‌های نهان‌نگاری مبتنی بر CNN را به طور خلاصه صورت‌بندی کرده‌است.

¹ Wang et al.

² Style Image

³ Prep-network

⁴ Hiding Network

⁵ Reveal Network

⁶ Baluja et al.

نه تنها نهان‌نگاری عکس، بلکه نهان‌نگاری فیلم نیز با استفاده از CNN مورد آزمایش قرار گرفته است. معمولاً از لایه‌های کانولوشنی دو بعدی برای عکس‌ها استفاده می‌شود در حالی که لایه‌های کانولوشنی سه بعدی برای فیلم‌ها مورد استفاده قرار می‌گیرد. پوشش موقت و قاب‌های ویدیوی مخفی به عنوان ورودی به VStegNET مبتنی بر شبکه رمزگذار خودکار [۳۴] برای تولید ویدیوی حامل^۱ داده می‌شود. هر فریم از تصویر رویه با هر فریم از ویدیوی مخفی الحاق می‌شود تا ویدیوی حامل تولید شود. در انتها از یک معماری شبکه یکسان برای استخراج ویدیوی پنهان استفاده می‌شود.

جدول ۲- مختصری از روش‌های نهان‌نگاری مبتنی بر شبکه‌های عصبی کانولوشنی

روش	معماری	مجموعه داده‌ها	مزایا	معایب
[۳۱]	رمزگذار-رمزگشا	ImageNet	-پیام مخفی، عکس است	در هر حال، تصویر 64x64 است که بسیار کوچک است.
[۲۷]	U-Net	ImageNet	-پیام مخفی، عکس است. -از معماری کمینه و ابتدایی استفاده شده است.	-در هر حال، تصویر 64x64 است که بسیار کوچک است. -عکس‌های ورودی فقط متصل به هم هستند.
[۲۸]	CNN	ImageNet و Holiday	-پیام مخفی، عکس است. از معماری کمینه و ابتدایی استفاده شده است. -تابع جدید خطای Back Propagation برای سرعت بخشی به آموزش، معرفی شده است.	-در هر حال، تصویر 64x64 است که بسیار کوچک است. -عکس‌های ورودی فقط متصل به هم هستند.
[۲۹]	رمزگذار-رمزگشا با پایه VGG	COCO و Wikiart.org	-Domain-Knowledge لازم ندارد. -امنیت بالای عکس تولید شده به اطلاعات مخفی مرتبط نیست.	محاسبات با استفاده از تصویر اضافی، افزایش می‌یابند.
[۲۵] و [۲۶]	رمزگذار-رمزگشا با SCR	ImageNet	امنیت بالا و قدرتمند	-Loss استفاده شده، بهینه نیست. -نویز قابل مشاهده در قسمتهای سفید یا سیاه، قابل دیدن است.

نه تنها نهان‌نگاری تصویر، بلکه نهان‌نگاری ویدئو نیز با استفاده از CNN امتحان شده است. معمولاً از لایه‌های کانولوشن دو بعدی برای تصاویر استفاده می‌شود در حالی که لایه‌های کانولوشنی سه بعدی برای فیلم‌ها استفاده می‌شود. پوشش موقت و قاب‌های ویدیوی مخفی به عنوان ورودی به VStegNET مبتنی بر شبکه رمزگذار خودکار [۳۴] برای تولید ویدیوی حامل داده می‌شود. هر فریم از تصویر رویه با هر فریم از ویدیوی

^۱ container video

مخفی الحاق می‌شود تا ویدیوی حامل تولید شود. از یک معماری شبکه یکسان برای آشکارکردن ویدیوی مخفی پنهان استفاده می‌شود.

۲-۲- روش‌های نهان‌نگاری مبتنی بر شبکه‌های مولد تخصصی^۱

شبکه‌های مولد تخصصی، نوعی از CNNهای عمیق هستند که توسط گودفلو^۲ و همکاران [۳۵] در سال ۲۰۱۴ معرفی شدند. یک GAN از نظریه بازی برای آموزش یک مدل مولد با فرآیند خصمانه برای وظایف تولید تصویر استفاده می‌کند. دو شبکه مولد^۳ و شبکه ممیز^۴ برای ایجاد یک تصویر عالی در معماری GAN با یکدیگر رقابت می‌کنند. به مدل مولد داده وارد می‌شود و خروجی، تقریبی نزدیک به تصویر ورودی داده شده است. شبکه‌های تشخیص دهنده تصاویر تولید شده را به عنوان جعلی^۵ یا واقعی^۶ طبقه بندی می‌کند. دو شبکه به گونه‌ای آموزش داده شده‌اند که مدل مولد سعی می‌کند داده‌های ورودی را تا حد امکان با حداقل نویز تقلید کند. مدل ممیز برای کشف موثر تصاویر جعلی آموزش داده شده است. از آن زمان تاکنون تغییرات زیادی در GAN ارائه شده است که آن را برای کارهای تولید تصویر مصنوعی قدرتمندتر و مناسب‌تر کرده است. GANها به دلیل عملکرد خوبشان در زمینه تولید تصویر شناخته می‌شوند. نهان‌نگاری تصویر را می‌توان به عنوان یکی از وظایف تولید تصویر در نظر گرفت که در آن دو ورودی تصویر روبه و تصویر مخفی برای تولید یک تصویر پوششی خروجی داده می‌شود. روش‌های موجود برای نهان‌نگاری تصویر با استفاده از معماری GAN را می‌توان به پنج گروه دسته‌بندی کرد - یک مدل GAN مبتنی بر شبکه، معماری‌های مبتنی بر چرخه GAN، معماری فرستنده-گیرنده با استفاده از GAN، مدل بدون پوشش که در آن تصویر پوشش به جای اینکه به طور تصادفی تولید شود. به عنوان ورودی و یک مدل مبتنی بر آلیس، باب و حوا ارائه شده است. جزئیات مربوط به همه دسته‌ها و نحوه اجرای آنها در زیر آورده شده است.

به طور کلی، یک مدل GAN از دو جزء اصلی تشکیل شده است: مولد و ممیز. در زمینه نهان‌نگاری تصویر، شبکه جدیدی به نام نهان‌تحلیلگر^۷ در برخی از روش‌ها معرفی شده است. وظایف اصلی این سه جزء عبارتند از:

- یک مدل مولد، G، برای تولید تصاویر پوششی از تصویر روبه و پیام تصادفی.
- یک مدل تشخیص دهنده، D، برای طبقه‌بندی تصویر تولید شده از مولد به عنوان واقعی یا جعلی.

^۱ GAN-Based Networks (Generative Adversarial Networks)

^۲ Goodfellow

^۳ Generator

^۴ Discriminator

^۵ Fake

^۶ True

^۷ steganalyzer

- یک نهان‌تحلیلگر، S ، برای بررسی اینکه آیا تصویر ورودی دارای داده‌های مخفی محرمانه هست یا خیر.

سه مدل G ، D و S برای رقابت با یکدیگر برای تولید تصاویر واقعی نزدیک به تصویر رویه ارائه شده به عنوان ورودی ساخته شده‌اند. خطاهای D و S به همراه پارامتر آلفا بین $[0, 1]$ برای تولید تصویر واقعی و کیفیت آنها برای نهان‌تحلیل^۱ استفاده می‌شود. یکی از تفاوت‌های اصلی با GAN در اینجا این است که G به منظور به حداکثر رساندن نه تنها خطای D ، بلکه برای به حداکثر رساندن خطای ترکیب خطی طبقه‌بندی‌دهای D و S ، به روز می‌شود.

ولکونسکی و همکاران^۲ [۳۶] [۳۷]، مدل DCGAN [38] مبتنی بر SGAN3 معرفی کرده است که یک DCGAN ساده با سه ماژول G ، D و S است. مشابه [۳۶]، شی و همکاران^۴ [۲] و [۳۹] یک معماری GAN سه جزء را پیشنهاد کرده است. تنها تفاوت در معماری استفاده شده توسط این روش‌ها است. یک DCGAN در [۳۶] و [۳۷] استفاده می‌شود، در حالی که یک لایه چهار کسری کانولوشن و به دنبال آن یک لایه عملکرد با فعال‌ساز تانژانت هیپربولیک با پایه WGAN در [۲] و [۳۹] استفاده می‌شود. سه مدل به گونه‌ای با یکدیگر رقابت می‌کنند که مولد تصاویر پوششی را تولید می‌کند و ممیز، پیام مخفی را رمزگشایی و بازیابی می‌کند در حالی که نهان‌تحلیلگر به مولد در حال تولید احتمال، گوش می‌کند. شکل ۳ اصل کار SGAN را نشان می‌دهد. گاهی اوقات یک شبیه‌ساز تعبیه در محل نهان‌تحلیلگر استفاده می‌شود [۴۱]، [۴۲] و [۴۳]. یانگ و همکاران [۴۱] و [۴۲] یک نهان‌نگاری تصویر مبتنی بر GAN با سه ماژول مولد، شبیه‌ساز جاسازی شده^۵ و ممیز، ارائه کردند. نویسندگان این معماری را UT-GAN نامیده‌اند [۴۱]، زیرا از معماری مبتنی بر U-Net در شبکه مولد استفاده می‌کنند. مولد برای ایجاد نقشه احتمال استفاده می‌شود، P ، برای تصویر پوشش ورودی، C و U-Net به دلیل عملکرد موثر آن در تقسیم بندی پیکسلی در نظر گرفته می‌شود. U-Net مورد استفاده در اینجا همچنین دارای یک لایه کانولوشن در حال گسترش و انقباض با افزایش تعداد فیلتر برای کانولوشن و کاهش تعداد فیلتر برای لایه‌های دی-کانولوشن^۶ است. همچنین لایه‌های I از قسمت نمونه‌برداری پایین با لایه‌های $L-i$ از قسمت نمونه‌برداری بالا به هم متصل می‌شوند تا به فرآیند پس‌انتشار کمک کنند. نقشه احتمال، P ، از مولد و پیام تصادفی به عنوان ورودی به شبیه‌ساز تعبیه شده برای خروجی نقشه اصلاح، M داده می‌شود. شبیه‌ساز جاسازی شده از یک تابع تانژانت هیپربولیک دوگانه [۴۲] و یک تابع تانژانت هیپربولیک [۴۱] استفاده می‌کند، زیرا تانژانت هیپربولیک قابل تمایز است و از دست دادن گرادیان در طول انتشار به عقب

¹ Steganalysis

² Volkhonskiy et al.

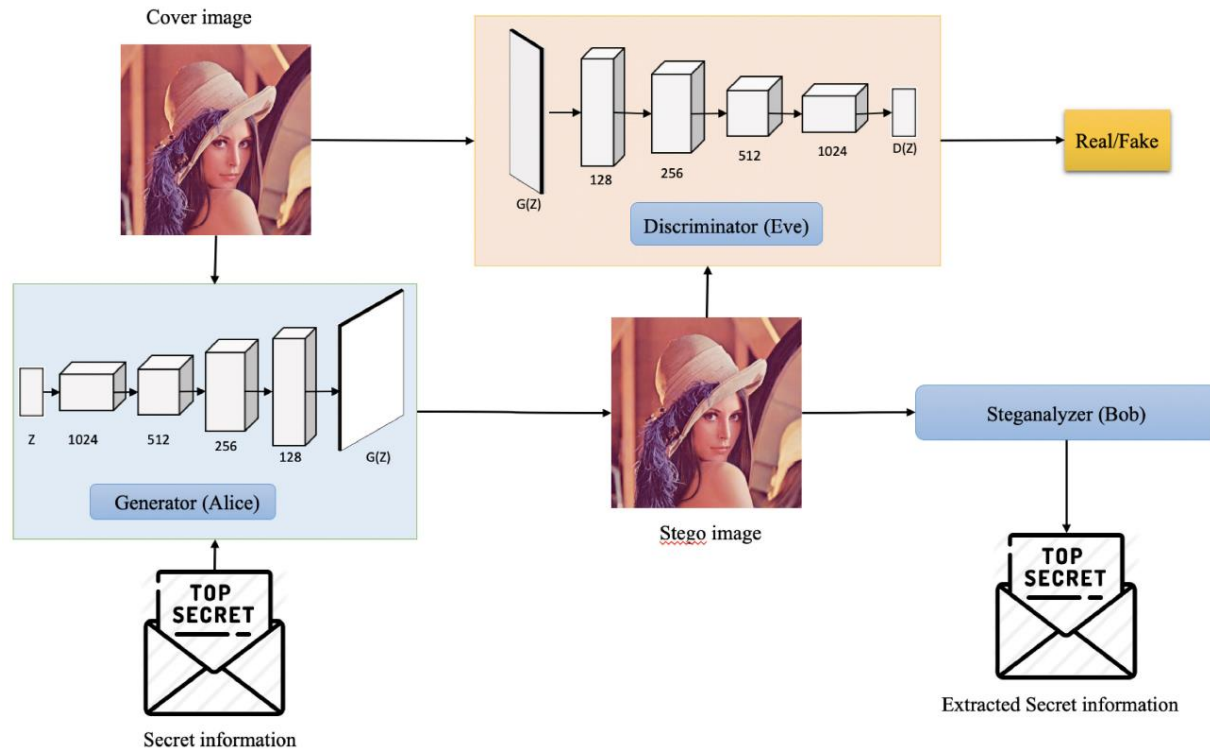
³ Steganographic GAN (SGAN)

⁴ Shi et al.

⁵ Embedding Simulator

⁶ deconvolution

را حفظ می‌کند. تشخیص‌دهنده به عنوان یک نهان‌تحلیلگر با چندین فیلتر بالا گذر برای پیش پردازش برای بهبود دقت تشخیص عمل می‌کند. دو گونه تلفات وجود دارد که در اینجا باید در نظر گرفته شود، تلفات ممیز و تلفات مولد. تلفات ممیز هنگام استفاده از آن در تمایز بین تصویر رویه و تصویر پوششی است در حالی که تلفات مولد بار هدف را نیز در برمی گیرد. بنابراین، تلفات مولد از دو تلفات تشکیل می‌شود: اتلاف



شکل ۳ سرریز کاری عمومی SGAN و [۴۰]. مولد تصویر رویه و اطلاعات مخفی (معمولاً متن) را به عنوان ورودی برای تولید تصویر پوششی می‌گیرد. ممیز با مولد رقابت می‌کند تا تصویر تولید شده را به عنوان واقعی یا جعلی طبقه‌بندی کند. یک نهان‌تحلیلگر اضافه می‌شود تا اطلاعات مخفی را از تصویر پوششی تولید شده، استخراج کند.

آنتروپی برای تضمین ظرفیت جاسازی و اتلاف تخاصم برای افزایش نرخ ضدتشخیص^۱.

چارچوب یادگیری خودکار اعوجاج نهان‌نگاری با GAN (به اختصار ASDL-GAN) توسط تانگ و همکاران در [۴۳] معرفی شد. در این معماری، از مولد برای یادگیری احتمالات هر پیکسل از تصویر پوششی ورودی استفاده می‌شود و نویسندگان پیشنهاد کرده‌اند که از یک تابع فعال‌سازی جدید شبیه‌ساز تعبیه سه‌گانه (TES) برای تولید تصاویر پوششی از احتمالات تولید شده استفاده شود. ممیز به تمایز بین تصاویر واقعی و جعلی کمک می‌کند.

معماری مبتنی بر XuNet برای ممیز D استفاده می‌شود. نوع معماری فرستنده-گیرنده راه دیگری برای نهان‌نگاری تصویر با استفاده از GAN است [۳۳]، [۴۳] و [۴۴]. Stegano GAN، یک مدل GAN با شبکه رمزگذاری

^۱ anti-detectability rate

برای ایجاد تصاویر پوششی و رمزگشایی شبکه برای دریافت پیام مخفی از تصویر پوششی توسط Zhang و همکارانش در [۴۵] پیشنهاد شده است. پیام مخفی با استفاده از رمزگذار در تصویر رویه تعبیه شده است و پیام مخفی توسط رمزگشا بازیابی می‌شود و از یک منتقد برای ارزیابی کیفیت تصاویر تولید شده استفاده می‌شود.

سه نوع رمزگذار پایه، باقیمانده و متراکم نیز برای مقابله با ظرفیت بار مختلف مورد بحث قرار گرفته است. سه مدل یک مدل پایه و دو مدل پیشرفته توسط چن و همکاران در [۴۳] معرفی شده‌اند. هر سه مدل برای انجام نهان‌نگاری با استفاده از شبکه رمزگذاری و بازیابی ایمن تصویر مخفی با استفاده از شبکه رمزگشایی در نظر گرفته شده‌اند که در آن مدل‌های بهبود یافته ایمن‌تر و قوی‌تر هستند. مدل پایه از شبکه رمزگذاری برای مخفی کردن تصویر مخفی مقیاس خاکستری در کانال رنگی B تصویر رویه استفاده می‌کند تا تصویر پوششی را تولید کند. به همین ترتیب، شبکه رمزگشایی تصویر مخفی را از کانال رنگی B تصویر پوششی نشان می‌دهد. دلیل استفاده از کانال B این است که تاثیر رنگ آبی بر چشم انسان کمتر از قرمز و سبز است. یک شبکه نهان‌واری و شبکه حمله در مدل پایه اضافه شده است تا در مدل‌های پیشرفته ایمن‌تر و قوی‌تر شود. XuNet با یک لایه هرمی فضایی Pooling در شبکه نهان‌تحلیلگر استفاده می‌شود. مدل مولد به عنوان فرستنده و تشخیص دهنده به عنوان گیرنده توسط Naito و همکاران [۴۴] استفاده می‌شود. مولد در انتهای ارسال، تصاویر پوششی را تولید می‌کند و ممیز، تصاویری را که واقعی نیستند حذف می‌کند. در انتهای دریافت، ممیز فرستنده تصویر پوششی دریافت شده را به فرستنده واقعی یا شخص ثالث طبقه بندی می‌کند تا از اتلاف زمان تولید کننده بر روی تصاویر پوششی شخص ثالث جلوگیری کند. مولد در انتهای دریافت کننده تصویر پوششی را رمزگشایی می‌کند و داده‌های مخفی را آشکار می‌کند. هم فرستنده و هم گیرنده از یک مولد آموزش دیده و ممیز برای نتایج ثابت استفاده می‌کنند.

HIDDeN یک روش مبتنی بر GAN با چهار جزء اصلی رمزگذار، رمزگشا، تشخیص دهنده و یک لایه نویز است که توسط Zhu و همکارانش پیشنهاد شده است. [۴۶]. رمزگذار تصویر رویه و پیام مخفی را به عنوان ورودی می‌گیرد و یک تصویر رمزگذاری شده ایجاد می‌کند که برای تولید تصویر نویز به لایه نویز داده می‌شود. رمزگشا پیام مخفی را از تصویر نویزی شده رمزگشایی می‌کند و تشخیص دهنده کمک می‌کند تا احتمال رمزنگاری یا عدم رمزنگاری تصویر داده شده را بدهد. که و همکاران^۱ [۴۷] معماری کمی متفاوت را با استفاده از نهان‌نگاری تولیدی با اصل کریشهوف^۲ (GSK) پیشنهاد کرده‌اند. به جای اصلاح تصویر رویه و تبدیل آن به تصویر پوششی، یک تصویر پوششی جدید با پیام مخفی تولید می‌شود. کلید استخراج و تصویر پوششی تولید شده برای رمزگشایی به ممیز داده می‌شود. بدون کلید استخراج، ممیز نمی‌تواند تصویر مخفی را رمزگشایی کند. مولد برای عموم در دسترس است، فرستنده ابتدا تصویر رویه را می‌دهد تا کلید استخراج را دریافت کند. سپس تصویر رویه و کلید استخراج برای گیرنده ارسال می‌شود. در مواردی که تصویر رویه یا

^۱ Ke et al.

^۲ Kerckhoffs' principle

کلید استخراج به تنهایی ارسال می‌شود، تنها نویز به عنوان خروجی داده می‌شود. تفکیک کننده تنها زمانی می‌تواند پیام مخفی را خروجی دهد که هم تصویر و هم کلید استخراج وجود داشته باشد.

روش CycleGAN [۴۸] برای نهان‌نگاری تصویر که در آن تصویر ورودی داده می‌شود و خروجی مشابه تصویر ورودی داده شده اما با اطلاعات پنهان با استفاده از آموزش تخصصی تولید می‌شود، مناسب است. تصویر ورودی ابتدا به یک تصویر دامنه هدف تبدیل می‌شود و سپس به تصویر منبع باز می‌گردد و نیاز به تصویر خروجی را از بین می‌برد. روش اصلی cycleGAN کمی اصلاح شده است تا کاملاً متناسب با روش نهان‌نگاری تصویر برای پنهان کردن پیام مخفی در داخل تصویر رویه [۴۹] [۵۱] و [۵۲] باشد.

ابتدا تصویر رویه RGB به تصویر در مقیاس خاکستری تبدیل می‌شود و سپس نواحی لوما^۱ در تصویر استخراج می‌شوند [۴۹]. پیام مخفی در بیت LSB فیلد لوما تعبیه شده است. سپس مولد تصویر جدید را با پیام مخفی تعبیه شده ایجاد می‌کند. تصویر تولید شده و تصویر رویه به ممیز داده می‌شود تا آن را جعلی یا واقعی طبقه بندی کند. یک حذف کننده بین دو مولد چرخه GAN برای کاهش و فیلتر کردن دامنه کم، پیام‌های فرکانس بالا و نویزهایی که تشخیص دهنده قادر به دیدن آنها نیست، معرفی شده است [۵۲]. یک چرخه GAN برای نهان‌نگاری و نهان‌واری در یک ارتباط مخفی برای جلوگیری از نقض امنیت و حفظ حریم خصوصی در IoT توسط منگ^۲ و همکاران در [۵۲] استفاده می‌شود. ACGAN [۵۳] می‌تواند تصاویر واقعی را برای یک برچسب مشخص تولید کند و همچنین برچسب تصاویر تولید شده را تشخیص دهد.

سه مرحله برای مخفی کردن اطلاعات تولید شده توسط مدل و استخراج اطلاعات پنهان [۵۴] و [۵۵] دنبال می‌شود. ابتدا یک فرهنگ لغت تقسیم بندی کلمات و پایگاه داده تصویر ایجاد می‌شود. یک مدل مخفی مولد به نام Stego-ACGAN توسعه یافته است. در نهایت یک الگوریتم مخفی و استخراج برای مخفی کردن و بازیابی اطلاعات طراحی شده است. یک پایگاه داده تنظیم شده است که کلمه بخش بندی را با عنوان برچسب و تصاویر مربوطه، جفت می‌کند. مانند هر ACGAN دیگری، stego-ACGAN نیز دارای سه شبکه عصبی است. مولد، تشخیص دهنده و طبقه‌بند کمکی. پس از آموزش stego-ACGAN، از یک کانال مخفی برای به اشتراک گذاشتن پارامترهای مدل و پایگاه داده‌های تصویر و بخش کلمه ساخته شده با طرف مقابل استفاده می‌شود. اطلاعات مخفی به بخش‌های تقسیم شده و یک کد باینری تولید می‌شود. سپس به هر کد یک برچسب زده می‌شود و حالت دنباله‌ای از تصاویر را برای بخش‌ها ایجاد می‌کند. دنباله‌های تصویر و نویز ورودی به مدل مولد داده می‌شود تا تصاویر پوششی تولید شود. در انتهای دریافت، نویز حذف می‌شود و دنباله‌های تصویر استخراج می‌شوند. سپس توالی تصویر به طبقه بندی کننده کمکی داده می‌شود تا اطلاعات محرمانه را بدست آورد. مشابه [۵۴] و [۵۵]، دوان^۳ و همکاران در [۵۶] یک روش نهان‌نگاری بدون پوشش با استفاده از مدل مولد پیشنهاد کرده‌اند. به جای انتقال تصویر مخفی به این شکل، با استفاده از

¹ Luma

² Meng

³ Duan

مدل مولد، یک تصویر معنی-عادی^۱ جدید که کاملاً به تصویر دیگر مربوط نمی‌شود، تولید می‌شود. در انتهای دریافت، تصویر معنی-عادی ارسال شده برای تولید تصویر مخفی تغذیه می‌شود. یک WGAN [۵۷] به عنوان مدل مولد توسط لی و همکاران^۲ در [۵۸] استفاده شده است. چارچوبی که در آن یک تصویر بافتی توسط یک مدل تولیدی تولید می‌شود و به عنوان یک تصویر پوششی عمل می‌کند، پیشنهاد شده است. سپس این تصویر رویه و تصویر مخفی به شبکه پنهان کننده داده می‌شود تا تصویر مخفی در داخل تصویر کاور مخفی شود. بنابراین تصویر نهایی یک تصویر مبتنی بر بافت است که با اطلاعات مخفی پنهان شده است. یک روش مبتنی بر یادگیری تخصصی با سه جزء آلیس، باب و حوا توسط هیز و همکاران^۳ [۴۰] پیشنهاد شده است.

جدول ۳ خلاصه‌ای از جزئیات روش‌های نهان‌نگاری مبتنی بر GAN.

روش	معماری	مجموعه داده	مزایا	معایب
[۳۹]	Alic, Bob and Eve	BOSSbase and celebA	-دانشتن دامنه در جاسازی الزام ندارد	- پیام بیشتر از تصاویر مورد مصرف قرار گرفته - جستجوی نقطه برای انتخاب شمای جاسازی شده زمان‌بر است.
[۳۷]	Info-GAN	MNIST and celebA	-امنیت اضافی بوسیله اضافه کردن یک کلید استخراج به همراه تصویر پوشش برای مولد، گسترش یافته است.	-پیام به وسیله یک مولد کلید استخراج، رمزگذاری شده‌است. کلید استخراج و مولد آموزش دیده به صورت عمومی در دسترس‌اند که آنها را مستعد حملات می‌کند.
[۲]	DCGAN	BOSSbase and celebA	-از فرمول بندی نظریه بازی‌ها بهره برده است	-از سه جزء بهره می‌برد. نهان‌تحلیگر استفاده شده خودش به تنهایی می‌تواند احتمال سرریز محاسباتی را افزایش دهد.
[۵۴]	WGAN	celebA	-تصویر همان پیام مخفی است.	-اگر مدل مولد خوب کار نکند، تصویر مخفی از بین خواهد رفت.
[۵۳]	ACGAN	MNIST	-امنیت بالا و توانمند با ظرفیت مخفی سازی گسترش یافته.	-طراحی پیچیده برای پنهان کردن اطلاعات مخفی. اولین تصاویر به ازای اطلاعات مخفی تولید می‌شوند و نویزهای اضافی باید اضافه شوند. -یک کانال امن برای انتقال اطلاعات لغوی پایگاه داده و سپس تصاویر لازم است.
[۴۷]	Cycle GAN	USC-SIPI	-کارایی و امنیت بهتر نسبت به روش‌های LSB.	-اطلاعات متنی مخفی می‌شوند. سوال اصلی اینجا این است که چگونه اطلاعات پنهان را از تصاویر استخراج می‌کنند. روش پیشنهادی کامل نیست. -گم شدن همگرایی GAN ممکن است منجر به فروپاشی شود.

¹ meaning-normal

² Li et al.

³ Hayes et al.

[۵۶]	DCGAN	DTD and COCO2017	-تصویر مخفی می‌تواند هر عرض و طولی داشته باشد.	-- صرفاً تصویر مبتنی بر بافت ساخته شده است.
[۵۰]	Cycle GAN	ImageNet and Photo2Monet	-امنیت بخشی به اطلاعات پزشکی برای گسترش حریم خصوصی اطلاعات ذخیره شده پزشکی.	- Cycle GAN برای پوشاندن داده‌های واقعی در IoT به کار رفته است. استفاده از نهان‌نگاری در جایگاه رمزگذاری توضیح داده نشده است.
[۴۳]	GAN	celebA	-تولید تعداد نامتناهی از تصاویر رویه امکانپذیر است.	-مولد و ممیز مشترک. کاربر باید مقدار کم را به جای تصویر رویه انتخاب کند. - تصاویر تولید شده، طبیعی نیستند.
[۳۸]	DCGAN	celebA	-تولید تصاویر واقعی تر. -امنیت بسیار بالا	- نهان‌تحلیگر احتمال بیشتری نسبت به اطلاعات پنهان دارد. -توضیحی درباره اینکه چگونه پیام پنهان را از تصویر رویه درمی‌آوریم، نیست.
[۴۲]	ASDL GAN	BOSSBase	-تابع فعال‌ساز جدید برای تولید تصاویر پوششی.	-GAN-ASD_ هنوز پایین تر از بهترین الگوریتم‌های نهان‌نگار است.

اصل کار این است که حوا، بین آلیس و باب استراق سمع می‌کند تا بررسی کند که آیا پیام مخفی در کانال ارتباطی بین آنها وجود دارد یا خیر. نویسندگان در [۴۰] از شبکه‌های عصبی برای آموزش هر سه مؤلفه استفاده کرده‌اند. برای سناریوی نهان‌نگاری، آلیس برای ایجاد تصویر نهان‌نگاری آموزش دیده است در حالی که باب پیام مخفی را از تصاویر پوششی بازیابی می‌کند. حوا با دادن احتمال اینکه تصویر داده شده یک تصویر پوششی باشد، به باب کمک می‌کند.

مدلی با چهار قسمت آلیس، باب، دیو^۱ و حوا توسط وانگ و همکاران در [۵۹] ارائه شده است. از آنجایی که یک مدل مولد بدون نظارت است، نویسندگان این مدل GAN را خود نظارتی نهان‌نگار (SSteGAN) نامگذاری کرده‌اند. مانند هر پارادایم امنیت ارتباطی دیگری، آلیس و باب سعی می‌کنند مخفیانه با هم ارتباط برقرار کنند در حالی که حوا از کانال ارتباطی استراق سمع می‌کند. به این ترتیب، آلیس به‌عنوان مولد، باب به‌عنوان رمزگشا، حوا به‌عنوان نهان‌تحلیگر برای طبقه‌بندی اینکه آیا تصویر داده‌شده عادی است یا تصویر پوششی، و دیو در اینجا به‌عنوان ممیزی عمل می‌کند که فعالانه با آلیس رقابت می‌کند. همراه با پیام مخفی، نویز ورودی نیز به عنوان ورودی به آلیس داده می‌شود تا در صورت دوبار دادن یک پیام مخفی، از ایجاد تصاویر یکسان جلوگیری شود. این امر هرگونه سوءظن ایجاد شده را منحرف می‌کند و امنیت مدل را افزایش می‌دهد. حوا و دیو، هم تصویر واقعی و هم تصویر تولید شده را می‌گیرند. در حالی که حوا به تشخیص پوششی و تصویر رویه کمک می‌کند، دیو در طبقه‌بندی تصویر به عنوان واقعی یا جعلی کمک می‌کند. خلاصه‌ای دقیق از روش‌های بررسی شده تحت عنوان روش مبتنی بر GAN در جدول ۳ آورده شده است.

¹ Dev

۳- مجموعه داده‌های مورد استفاده

یک مجموعه داده به نام BOSSBase وجود دارد که به طور خاص برای مقابله با مشکلات نهان‌نگاری ایجاد شده است. برای ارزیابی بیشتر عملکرد الگوریتم‌ها، برخی از مجموعه داده‌های موجود، که برای اهداف دیگری از جمله تشخیص شیء و تشخیص چهره استفاده می‌شوند، دوباره مدل‌سازی شده‌اند تا برای اهداف آزمایش‌های ما مناسب باشند. توضیح دقیق در مورد مجموعه داده‌ها در جدول ۴ توضیح داده شده است.

۳-۱- BOSSBASE

مسابقه شکستن سیستم نهان‌نگاری ما^۱ (BOSS) [۶۰] اولین چالش علمی است که برای تبدیل نهان‌نگاری تصویر از یک موضوع تحقیقاتی به یک کاربرد عملی انجام شده است. هدف اصلی این مسابقه، توسعه یک روش نهان‌کاوی بهتر بود که بتواند تصاویر نهان‌نگاری ایجاد شده توسط الگوریتم HUGO (به شدت غیرقابل شناسایی) را بشکند [۶۱]. مجموعه داده شامل یک مجموعه آموزشی و مجموعه تست همراه با الگوریتم HUGO است که می‌توان از آن استفاده کرد.

برای ایجاد تصاویر نهان‌نگاری مجموعه داده آموزشی شامل ۱۰۰۰۰ تصویر رویه خاکستری با ابعاد ۵۱۲ * ۵۱۲ است. مجموعه آزمایشی متشکل از ۱۰۰۰ تصویر در مقیاس خاکستری با ابعاد ۵۱۲ * ۵۱۲ است. گزینه‌ای برای دانلود مجموعه داده‌ها با تصاویر نهان‌نگاری صرفاً به منظور استگانالیز وجود دارد. ابتدا تصاویر خام با استفاده از ۷ دوربین مختلف گرفته شده و به تصاویر PGM تبدیل می‌شوند. لینک‌های دانلود تصاویر خام، تصاویر PGM، اسکرپیت مورد استفاده برای تبدیل تصاویر خام به PGM، داده‌های EXIF تصاویر خام را می‌توانید در وب سایت رسمی پیدا کنید.

^۱ Break Our Steganographic System (BOSS)

جدول ۴ جزئیات اطلاعات مجموعه داده‌های مورد استفاده در ادبیات نهان‌نگاری.

Dataset	Number of samples	Images format	Image Size	Purpose
BOSSBase	9074 training and 1000 testing	tiff	512 × 512 × 1	Steganography and steganalysis
CelebA	More than 200K	jpg		Face attribute recognition, face detection, landmark (or facial part) localization, and face editing and synthesis
ImageNet	More than 14M	Arbitrary	Arbitrary	Computer Vision
MNIST Handwritten Digits	60,000 training and 10,000 testing	idx	28 × 28 × 1	Image processing and computer vision
COCO	330K	jpg	640×640×3	Object detection, segmentation and image captioning
Div2K	800 training, 100 validation and 100 testing	png	1020×678×3	Single Image Super-Resolution
SZUBase	40000	-	512×512×1	Steganography and steganalysis
USC-SIPI	170	tiff	256×256×1, 512×512×1, or 1024×1024×1	Image processing, image analysis, and machine vision
DTD	5640	jpg	300×300×3 and 640×640×3	Textural analysis
LFW	13233	jpg	150×150×3	Face verification
Pascal VOC	11530	jpg	500×300×3	Object detection and classification

۲-۳- CELEBA

مجموعه داده CelebFaces Attributes در مقیاس بزرگ، همچنین به عنوان مجموعه داده CelebA [۶۲] شناخته می‌شود، مجموعه داده وسیعی با بیش از ۲۰۰ هزار تصویر است که می‌تواند برای تشخیص چهره، تشخیص چهره، محلی‌سازی چهره و سایر عملیات مرتبط با چهره استفاده شود. مجموعه داده شامل تصاویری از منابع، مکان‌ها، پس‌زمینه و ژست‌های مختلف است و برای نهان‌نگاری نیز مناسب است. احتمال استفاده از تصویر عکس/چهره به عنوان پوششی برای مخفی کردن تصاویر مخفی بسیار زیاد است. همراه با تصاویر، ۴۰ حاشیه نویسی مختلف مانند با/بدون عینک، احساسات، مدل مو، لوازم جانبی دیگر مانند کلاه، موجود است.

۳-۳- IMAGENet

همچنین، مجموعه داده بسیار بزرگی است [۶۳] که حاوی تصاویری از سلسله مراتب WordNet است که هر گره حاوی بیش از ۵۰۰ تا ۱۰۰۰ تصویر است. ImageNet هیچ حق چاپی برای تصویر ندارد و فقط شامل پیوندها یا تصاویر کوچک به تصویر اصلی است. مجموعه داده شامل تصاویر با اندازه‌های مختلف است. بر اساس نیاز، تعداد تصاویر، کلاس‌هایی که به آنها تعلق دارند، پس زمینه و اندازه تصویر را می‌توان از طیف گسترده موجود انتخاب کرد.

۴-۳- ارقام دست نویس MNIST

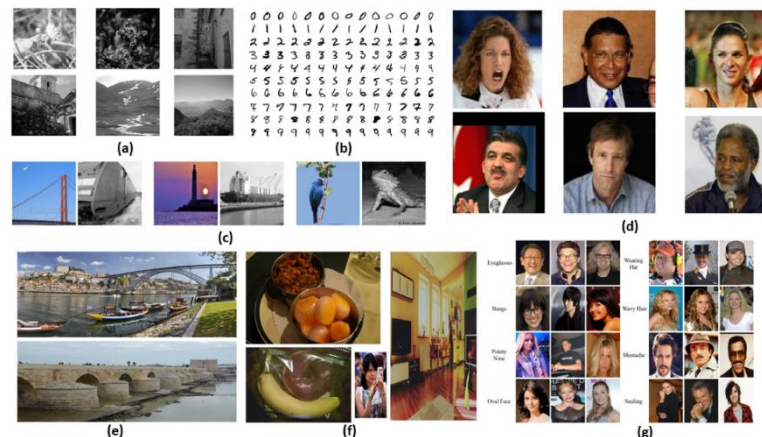
پایگاه داده موسسه ملی استاندارد و فناوری اصلاح شده^۱ (MNIST) [۶۴] مجموعه دیگری است که می‌تواند برای برنامه‌های مختلف بینایی کامپیوتری و پردازش تصویر استفاده شود. مجموعه داده دست‌نویس MNIST شامل یک مجموعه آموزشی و آزمایشی با تصاویری از ارقام دست‌نویس ۰ تا ۹ است. تصاویر در این مجموعه داده‌ها عادی، سیاه و سفید با ابعاد 28×28 پیکسل هستند. مجموعه آموزشی شامل ۶۰۰۰۰ تصویر و مجموعه تست شامل ۱۰۰۰۰ تصویر می‌باشد.

۵-۳- COCO

مجموعه داده‌های^۲ (COCO) [۶۵] عمدتاً برای اهداف تشخیص اشیاء، تقسیم بندی و شرح تصاویر ایجاد شده است. این یک مجموعه داده عظیم با تصاویر از ۸۰ دسته شیء است. هر کلاس شامل حداقل ۵ تصویر است. این مجموعه داده همراه با حاشیه نویسی کلاس و حاشیه نویسی تقسیم بندی می‌باشد. هیچ تقسیم آموزشی و آزمایشی از پیش تعریف شده وجود ندارد. تقسیم مجموعه داده را می‌توان بر اساس موضوع تحقیق و راحتی کاربر انجام داد.

۶-۳- دیگر مجموعه‌ها

مجموعه داده‌های دیگری که می‌توانند برای نهان‌نگاری و نهان‌واری تصویر استفاده شوند عبارتند از Div2K،



شکل ۴- تصاویر نمونه از مجموعه داده‌های (a) BOSSBase (b) MNIST (c) ImageNet (d) CelebA (e) Div2K (f) COCO and (g) CelebA.

Pascal VOC، Div2K و LFW، DTD، USC-SIPI، SZUBase [۶۶] مجموعه داده‌ای است که معمولاً برای وضوح تصویر تکی استفاده می‌شود که در چالش NTIRE 2017 در وضوح تصویر معرفی شده است. در مجموع دارای ۱۰۰۰ تصویر است که به ۸۰۰ تصویر برای آموزش، ۱۰۰ تصویر برای اعتبار سنجی و ۱۰۰ تصویر برای آزمایش

¹ Modified National Institute of Standards and Technology

² Common Objects in Context

تقسیم شده است. تصاویر با وضوح بالا (۱۰۲۴*۶۷۸) و سه درجه وضوح پایین را می‌توان در این مجموعه داده یافت. SZUBase مجموعه داده‌ای است که در [۶۷] با ۴۰۰۰۰ تصویر در مقیاس خاکستری با اندازه ۵۱۲*۵۱۲ جمع آوری شده است. USC-SIPi [۶۸] دارای انواع تصاویر با وضوح و اندازه‌های مختلف برای پردازش تصویر، تجزیه و تحلیل تصویر و اهداف بینایی ماشین است. مجموعه داده‌های بافت‌های قابل توصیف (DTD) [۶۹] برای تجزیه و تحلیل اجزای بافتی حاوی ۵۶۴۰ تصویر jpg در دو اندازه - ۳۰۰*۳۰۰ و ۶۴۰*۶۴۰ استفاده می‌شود. مجموعه داده LFW برای کارهای تشخیص چهره و تایید استفاده می‌شود [۷۰]. PASCAL Visual Object Classes (VOC) چالشی است که برای تشخیص و طبقه بندی اشیاء انجام می‌شود [۷۱]. شکل ۴ زیر برخی از تصاویر گردآوری شده از مجموعه داده توصیف شده را نشان می‌دهد.

۴- ارزیابی

معیارهای ارزیابی برای اندازه گیری نامرئی بودن، امنیت، استحکام و ظرفیت روش‌های پیشنهادی استفاده می‌شود. متداولترین معیارهای مورد استفاده و عملکردی که آنها اندازه گیری می‌کنند، در زیر توضیح داده شده است.

۴-۱- راه‌اندازی تجربی

مدل‌های نهان‌نگاری تصویر عمدتاً از دو ورودی و دو تصویر خروجی استفاده می‌کنند. علاوه بر این، مجموعه داده‌های مورد استفاده معمولاً بسیار بزرگ هستند. یک کامپیوتر مبتنی بر GPU با یک کارت گرافیک قدرتمند برای انجام آموزش و تست مورد نیاز است. سپس مدل آموزش دیده را می‌توان در یک CPU یا رایانه‌های مستقل دستی، مستقر کرد. همه آزمایش‌ها با استفاده از نسخه‌های پایتون ۳/۵ و بالاتر با پایتورچ [۲۷]، [۳۲]، [۵۶] و [۲۸] یا متلب [۱۶]، [۱۷]، [۴۹] و [۱۰] انجام می‌شوند. یکی دیگر از کتابخانه‌های محبوب مورد استفاده، تنسورفلو^۱ [۴۰]، [۵۲]، [۵۵] و [۳۹] است. کتابخانه چینر^۲ از پایتون نیز استفاده می‌شود [۴۴].

۴-۲- معیارهای مورد استفاده

۴-۲-۱- PSNR

نسبت سیگنال به نویز پیک (PSNR) برای تعیین کیفیت، استحکام و نامرئی بودن روش نهان‌نگاری پیشنهادی استفاده می‌شود. PSNR نسبت بین نمایش حداکثر کیفیت تصویر رویه و تصویر پوششی است. در مورد نهان‌واری، نسبت اندازه‌گیری حداکثر کیفیت تصویر مخفی اصلی و تصویر مخفی استخراج شده

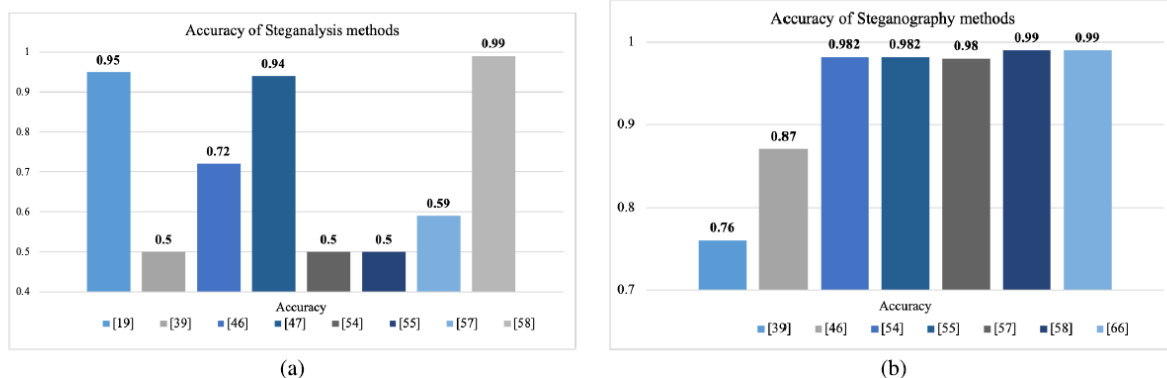
¹ TensorFlow

² Chainer

است. PSNR برای اندازه گیری خطای پیک روش پیشنهادی استفاده می شود. ارزش PSNR باید زیاد باشد؛ به این معنا که کیفیت تصویر پوششی بازسازی شده خوب است. میانگین مربعات خطا (MSE) معیار دیگری برای اندازه گیری کیفیت تصویر پوششی بازسازی شده است. MSE خطای مجذور تجمعی بین تصویر پوششی و تصویر رویه اصلی است. برای تصاویر با کیفیت بهتر، PSNR باید زیاد باشد در حالی که مقدار MSE باید کم باشد که نشان دهنده کم بودن خطا است. فرمول های محاسبه MSE و PSNR در زیر آورده شده است.

$$MSE = \frac{\sum_{M,N} [I_1(m,n) - I_2(m,n)]^2}{M*N} \quad (1)$$

که در آن، M و N به ترتیب تعداد سطرها و ستون ها در تصویر ورودی است. پس از محاسبه MSE، مقدار آن در محاسبه PSNR استفاده می شود.



شکل ۵ - مقایسه دقت نهان تحلیلگر میان روش های مختلف

$$PSNR = 10 * \log_{10} \frac{(R^2)}{MSE} \quad (2)$$

که در آن R نوسان در تصویر ورودی است. به عنوان مثال، اگر تصویر ورودی دارای یک نوع داده نقطه ممیز شناور با دقت دوگانه^۱ باشد، آنگاه R برابر ۱ است. اگر نوع داده عدد صحیح بدون علامت ۸ بیتی داشته

Method	PSNR	MSE	SSIM
[16]	62.53	-	-
[17]	56.95	0.13	-
[31]	-	-	6.4* and 3.6**
[27]	40.47* and 40.66**	-	0.9794* and 0.9842**
[47]	64.7	-	-
[56]	-	3.54 and 11.46**	-
[10]	32.09	-	-

جدول ۵ مقادیر PSNR و MSE روش ها. مقادیر نشان دهنده تصاویر پوششی برای تصاویر رویه* و مخفی** هستند.

^۱ double-precision floating-point

باشد، R برابر ۲۵۵ است. جدول ۵ مقادیر PSNR، MSE و SSIM را خلاصه کرده که با روش‌های مختلف در مطالعه به دست آمده است.

۲-۲-۴- دقت

دقت/بیت دقت^۱ معیاری است که معمولاً برای اندازه‌گیری امنیت و استحکام روش‌ها استفاده می‌شود. دقت نهان‌نگاری به این صورت است که دقت روش پیشنهادی به درستی تشخیص تصویر به عنوان تصویر نهان‌نگاری است یا خیر. دقت با استفاده از چهار عبارت محاسبه می‌شود: مثبت واقعی (TP)، منفی واقعی

جدول ۶ ماتریس سردرگمی برای محاسبه دقت.

Actual	Predictions	
	True	False
True	True Positive	False Negative
False	False Positive	True Negative

(TN)، مثبت کاذب (FP) و منفی نادرست (FN). مثبت واقعی و منفی درست پیش‌بینی‌های صحیح مدل هستند در حالی که مثبت کاذب و منفی کاذب خطاهای مدل هستند. جدول ۶ برای محاسبه معیارها از ماتریس سردرگمی^۲ استفاده می‌شود.

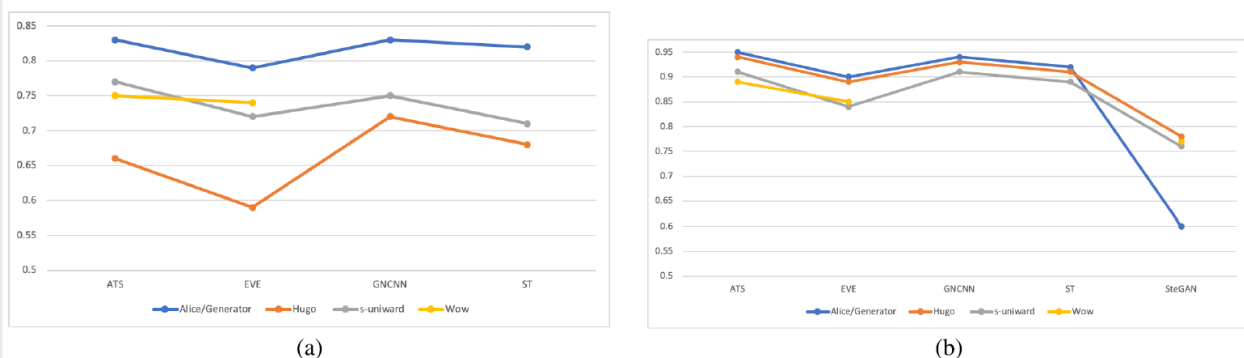
دقت به عنوان تعداد کل پیش‌بینی‌های صحیح انجام شده توسط مدل در برابر همه طبقه‌بندی‌ها محاسبه می‌شود. اگرچه دقت ساده و متداول است، اما زمانی که داده‌ها نامتعادل هستند، معیار خوبی نیست. این هزینه برای هر دو مثبت کاذب و منفی کاذب یکسان است. علاوه بر دقت، سایر معیارهای ارزیابی برای سنجش عملکرد واقعی مدل استفاده می‌شوند.

$$Accuracy = \frac{(TP+TN)}{(TP+FP+FN+TN)} \quad (3)$$

^۱ Accuracy bit

^۲ Confusion Matrix

برای یک روش نهان‌نگاری قوی‌تر، مقدار کمتر دقت نهان‌تحلیل، امنیت بالاتر را نشان می‌دهد. نرخ آشکارسازی^۱ عبارت دیگری است که برای محاسبه دقت آزمایش مدل نهان‌تحلیل استفاده می‌شود. برای تعیین محرمانه بودن روش پیشنهادی استفاده می‌شود. شکل ۵ دقت روش‌های نهان‌نگاری و نهان‌تحلیل مورد استفاده در مطالعه را نشان می‌دهد. مقایسه روش‌های دقت نهان‌تحلیل برای روش‌های مختلف برای تصاویر پوششی تولید شده با روش‌های رایج مانند [72] WOW، [61] HUGO، S-UNIWARD در شکل ۶ آورده



شکل ۶ مقایسه دقت بین روش‌های مختلف (الف) نهان‌تحلیل و (ب) نهان‌نگاری.

شده است.

۳-۲-۴ BPP

ظرفیت پنهان الگوریتم نهان‌نگاری با استفاده از بیت در هر پیکسل (BPP) محاسبه می‌شود. BPP تعداد بیت‌هایی را نشان می‌دهد که در هر پیکسل از تصویر پوششی برای تولید تصویر پوششی پنهان شده‌اند. برای ظرفیت پنهان بالاتر، ارزش BPP باید بالا باشد. برای محاسبه مقدار BPP از معادله زیر استفاده می‌شود.

$$BPP = \frac{L}{HWC} \quad (4)$$

که در آن L طول پیام پنهان، H ارتفاع، W عرض و C تعداد کانال‌های تصویر جلد است.

بسیاری از این معیارها با یکدیگر رقابت می‌کنند. مدل‌هایی که ظرفیت بالاتری دارند، معمولاً پنهان‌کاری را قربانی می‌کنند، زیرا بیشتر پنهان می‌شوند. اطلاعات موجود در تصاویر به طور طبیعی منجر به اعوجاج تصویر بزرگتر می‌شود. مدل‌هایی که در برابر نویز بسیار مقاوم هستند، معمولاً ظرفیت یا محرمانه بودن را قربانی می‌کنند، زیرا پیام باید به صورت اضافی در تصویر رمزگذاری شود. به نوعی، نهان‌نگاری و واترمارکینگ در دو سر طیفی از مشکلات قرار دارند که این محورهای مختلف را در اولویت قرار می‌دهند. نهان‌نگاری بر رازداری تاکید می‌کند در حالی که واترمارکینگ بر استحکام. در موارد مخفی کردن متن در یک تصویر، BPP استفاده شده یا ۰/۱ یا ۰/۴ است.

¹ Detection rate

۳-۴- مشاهده‌ها

مشاهده‌های انجام شده از نتایج گزارش شده توسط روش‌ها در اینجا مشخص شده است. اصولاً در هنگام بحث از مشاهدات، ظرفیت پنهان، عوامل امنیتی و استحکام در نظر گرفته می‌شود.

ظرفیت پنهان: به طور کلی ذکر می‌شود که ظرفیت پنهان سازی روش‌ها به ترتیب ذیل است. روش‌هایی که کمترین ظرفیت پنهان را دارند، روش‌های سنتی هستند که در آن متن شکل اولیه ارتباط مخفیانه است. پس از آن روش‌های مبتنی بر GAN است که در آن فقط پیام متنی به عنوان مخفی استفاده می‌شود. برخلاف روش‌های سنتی و مبتنی بر GAN، ظرفیت پنهان روش‌های مبتنی بر CNN به مراتب بهتر است و تقریباً ۱ است [۲۷]. اندازه تصویر مخفی مانند تصویر رویه است. حتی زمانی که یک تصویر در مقیاس خاکستری برای پنهان کردن استفاده می‌شود [۳۳]، اندازه تصاویر مخفی و رویه برابر است. از نظر ظرفیت پنهان، روش‌های مبتنی بر CNN به وضوح از سایر روش‌ها بهتر عمل می‌کنند.

امنیت و استحکام: امنیت با جاسازی و استحکام با استخراج تصویر مخفی مرتبط است. در مشاهده‌های ما، روش‌های مبتنی بر CNN و روش‌های مبتنی بر GAN امنیت بالاتری را به همراه دارند. شایان ذکر است که استخراج تصاویر مخفی در روش‌های یادگیری عمیق مستعد از دست دادن اطلاعات است. اما در روش‌های سنتی امنیت کمتر اما استحکام بالاست. معیار PSNR برای ارتباط بین امنیت و استحکام استفاده می‌شود. مقادیر از ۵ شروع می‌شود و می‌توان اشاره کرد که [۴۹] دارای بالاترین مقدار PSNR است که امنیت بالاتر روش مبتنی بر GAN را توضیح می‌دهد. بالاترین مقدار PSNR برابر ۶۴/۷ است که توسط چرخه نهان‌نگاری تصویر مبتنی بر GAN ارائه شده است.

از مشاهدات انجام شده، روش‌های یادگیری عمیق مبتنی بر GAN بهترین عملکرد را از نظر ظرفیت پنهان (۱) و PSNR (۶۴/۷) دارند [۴۹]. ممیزها به گونه‌ای آموزش دیده‌اند که بتوانند بر هرگونه حمله نهان‌تحلیل غلبه کنند و در مقایسه با روش‌های سنتی و مبتنی بر CNN، خاصیت ضد تشخیص بهتری دارند.

۵- چالش‌ها

در زیر چند چالش برای در نظر گرفتن مشکلات نهان‌نگاری تصویر آورده شده است.

در دسترس بودن داده‌ها - اگرچه استگنوگرافی تصویر یک یادگیری بدون نظارت است و هدف اصلی بازسازی تصویر است، هیچ مجموعه داده معیار مناسبی به جز BOSSBase [۶۰] وجود ندارد. ممکن است تعداد تصاویر در BOSSBase زیاد باشد، اما تصاویر در مقیاس خاکستری هستند که در قالب tiff در دسترس هستند. اکثر روش‌ها با پنهان کردن تصاویر RGB در داخل تصاویر رویه RGB سروکار دارند. یافتن یک مجموعه داده مناسب می‌تواند چالش برانگیز باشد. ImageNet متداول‌ترین مجموعه داده مورد استفاده است که یک اشکال عمده آن اندازه تصویر است. تصاویر در سایز ۶۴*۶۴ بسیار کوچک هستند.

همگرایی GAN - همگرایی یک اشکال بزرگ برای GAN است که در آن مدل بدون توجه به پارامترهای انتخاب شده، همگرا نمی‌شود. فروپاشی حالت نیز اغلب اتفاق می‌افتد زیرا مولد و ممیز به یکدیگر وابسته هستند.

مقایسه با روش‌های دیگر - معیارهای ارزیابی مورد استفاده در روش‌های مختلف متفاوت است و از این رو مقایسه روش پیشنهادی با روش‌های پیشرفته آسان نیست.

نهان‌نگاری بی‌درنگ^۱ - مدل‌های نهان‌نگاری بر روی مقادیر عظیمی از مجموعه داده‌ها آموزش داده می‌شوند، مانند [۲۷]، ۴۵۰۰۰ تصویر آموزشی استفاده شده است. با این حال، وقتی نوبت به نهان‌نگاری بی‌درنگ می‌رسد، کار دشواری می‌شود. اجرای مدل آموزش دیده برای انجام نهان‌نگاری و استگانالیز مستلزم انتقال تصویر استگو از طریق یک کانال غیرقابل اعتماد به انتهای گیرنده است. توانایی مدل آموزش دیده در برخورد با تصاویر زنده واقعی که ممکن است دارای نویز، کج شدن یا تار بودن باشد ثابت نشده است. اجرای مدل برای نهان‌نگاری بلادرنگ هنوز مشکوک است.

۶- بحث و کارهای آینده

GAN در مقایسه با روش‌های مبتنی بر CNN، و به طور خاص cycleGAN پرکاربردترین معماری است. مهمترین عاملی که باید در نظر گرفت این است که GAN یک مدل دو بخشی است که در آن یک مدل در انتهای فرستنده برای جاسازی اطلاعات مخفی استفاده می‌شود در حالی که مدل دوم در انتهای دریافت کننده برای استخراج اطلاعات مخفی استفاده می‌شود.

دو مدلی که تحت شرایط آموزشی یکسان آموزش داده می‌شوند تا کل فرآیند نهان‌نگاری تصویر به خوبی کار کند. از دست دادن یا نداشتن یک مدل ممکن است جاسازی/استخراج را تحت تأثیر قرار دهد زیرا آنها به یکدیگر متصل هستند. روش‌های مبتنی بر CNN از معماری رمزگشای خودکار UNet/Xu-Net برای جاسازی و استخراج استفاده می‌کنند. برخی از روش‌ها از رمزگذار برای جاسازی و از رمزگشا برای استخراج استفاده می‌کنند، در حالی که برخی از یک رمزگشای خودکار برای جاسازی و دیگری برای استخراج استفاده می‌کنند. اگرچه وابستگی متقابلی وجود دارد، اما مانند روش‌های GAN کاملاً مرتبط نیست.

عدم تعادل در یادگیری مولد-ممیز می‌تواند در جایی اتفاق بیفتد که مولد به طور موثر عمل می‌کند اما ممیز در تلاش است. اگرچه کارایی کلی تحت تأثیر قرار نمی‌گیرد، فرستنده یا گیرنده مستعد خطا هستند. با انتخاب دقیق پارامترها و اجتناب از تطبیق بیش از حد در طول تمرین می‌توان از این امر جلوگیری کرد.

^۱ Real-time steganography

قابلیت‌های امنیتی روش‌های GAN در مقایسه با معماری‌های CNN و روش‌های سنتی LSB بالاتر است. GAN‌ها اساساً در زمینه بازسازی تصویر مورد استفاده قرار می‌گیرند که آن را در مقایسه با convNets برای نهان‌نگاری تصویر، مناسب می‌کند.

در روش‌های یادگیری عمیق، اصل کار نهان‌نگاری تصویر استخراج ویژگی‌ها از روی جلد و تصویر مخفی و به هم پیوستن آنها برای ایجاد نتیجه نهایی نزدیک‌تر به تصویر رویه است. با این حال، کجا و چگونه پیکسل‌های تصویر مخفی جاسازی شده اند، نمی‌توان به وضوح درک کرد.

بدون آموزش مدل استخراج هم‌تا، ممکن است شکستن تصویر نهان‌نگاری دشوار باشد. این امنیت را افزایش می‌دهد، اما زمانی که مدل استخراج کار نمی‌کند یا خراب می‌شود، دشوار می‌شود.

برخی از معایب عمده عبارتند از زمان صرف شده برای آموزش، زمان محاسباتی در طول آزمایش و قابلیت‌های ذخیره سازی. مدل‌ها دو تصویر یا یک تصویر و یک پیام متنی را به عنوان ورودی به بیت می‌گیرند. ویژگی‌ها از هر دو ورودی استخراج می‌شوند که زمان محاسباتی را هم در تعبیه و هم در استخراج افزایش می‌دهد. تعداد پارامترها نیز در مقایسه با یک معماری معمولی حداقل دو برابر افزایش می‌یابد که به نوبه خود فضای ذخیره سازی مورد نیاز مدل را افزایش می‌دهد.

هنگامی که دیگران از تصاویر در مقیاس خاکستری استفاده می‌کردند، تصاویر مخفی RGB با روش‌های انگشت شماری استفاده می‌شوند. هنگام تبدیل تصویر در مقیاس خاکستری به تصویر RGB برای درک بهتر، ممکن است اطلاعات از دست برود. علاوه بر درک صحیح اطلاعات محرمانه، به تکنیک‌های بهبود تصویر نیز نیاز است.

در برخی از جنبه‌هایی که می‌توان برای کارهای آینده در نظر گرفت که در ادامه نیز برشمرده می‌شوند، استفاده از شبکه‌های محبوب U-Net، cycleGAN و استفاده از DCT و DWT در نظر گرفته شده‌اند و می‌توان در مورد سایر معماری‌های سفارشی‌شده کاوش بیشتری کرد. به عنوان مثال، شبکه‌های عصبی بازگشتی (RNN) به جای CNN‌ها را می‌توان بیشتر مورد بررسی قرار داد. یک WGAN سفارشی را می‌توان با انواع دیگر GAN جایگزین کرد.

- اکثر روش‌های نهان‌نگاری تصویر از متن یا تصویر در مقیاس خاکستری به عنوان اطلاعات محرمانه استفاده می‌کنند و نیاز به تحقیقات بیشتری در زمینه پنهان کردن تصویر در تصویر و ویدئو وجود دارد.

- آزمایش‌های مربوط به بهینه سازی پارامترها و کاهش ظرفیت‌های ذخیره سازی را می‌توان با استفاده از مجموعه داده‌های مختلف انجام داد.

- دوران محاسبات کوانتومی نگذشته است؛ تلاش‌های بیشتری در زمینه توسعه طرح‌ها بر روی تصاویر کوانتومی قابل بررسی است.

- برای بهره مندی از ترکیبی از روش‌ها، مجموعه‌ای از روش‌های یادگیری سنتی و عمیق را می‌توان بیشتر مورد مطالعه قرار داد.

- می‌توان تلاش‌ها را برای تشکیل یک مجموعه داده معیار حاوی تصاویر از دوربین‌های منبع مختلف، فرمت‌های تصویر، هدایت کرد. تلفیقی از تمام الگوریتم‌های ممکن نیز می‌تواند برای ایجاد تصاویر نهان‌نگاری انجام شود.

- بسیاری از روش‌ها ظرفیت پنهان سازی، امنیت و استحکام را به عنوان معیار عملکرد در نظر گرفته اند. با این حال، زمانی که انتقال از طریق کانال‌های غیرقابل اعتماد انجام می‌شود، احتمال حملات انسان در وسط وجود دارد. دستکاری تصویر پوششی نیز می‌تواند در حین انتقال اتفاق بیفتد. این حملات و عملکرد الگوریتم طراحی شده در برابر این حملات را می‌توان برای ارزیابی در کنار سایر معیارها در نظر گرفت.

۷- نتیجه‌گیری

نهان‌نگاری تصویر روشی است که در انتقال اطلاعات محرمانه با پنهان کردن آن در دید آشکار در داخل یک تصویر رویه استفاده می‌شود. روش‌های یادگیری عمیق به طور گسترده در هر زمینه‌ای مورد استفاده قرار می‌گیرد و در تحقیقات نهان‌نگاری مورد استفاده قرار گرفته است. بررسی تمام آثار مرتبط منجر به دسته بندی آنها به سه گروه شد. بسیاری از روش‌های نهان‌نگاری مبتنی بر سنتی از جایگزینی LSB و برخی از انواع آن استفاده می‌کنند. به غیر از LSB، روش‌های PVD، DCT و EMD معمولاً استفاده می‌شود. ظرفیت پنهان کردن روش‌های سنتی محدود است، زیرا بارگذاری بیش از حد تصویر جلد با بهره‌برداری از پیکسل‌های بیشتر برای پنهان کردن پیام مخفی ممکن است منجر به اعوجاج شود.

همچنین، ساختار رمزگذار-رمزگشا خودکار با VGG به عنوان پایه، U-Net و Xu-Net رایجترین معماری‌های مورد استفاده برای روش‌های نهان‌نگاری تصویر مبتنی بر CNN هستند. اخیراً، معماری GAN به دلیل توانایی آنها در مقابله با وظایف بازسازی تصویر توجه قابل توجهی را به خود جلب کرده است. نهان‌نگاری تصویر را می‌توان یکی از کارهای بازسازی تصویر در نظر گرفت که در آن تصویر جلد و اطلاعات مخفی به عنوان ورودی برای بازسازی یک تصویر نهان‌نگاری که از نظر شباهت به تصویر رویه نزدیک است، گرفته می‌شود.

هیچ مجموعه داده تصویر معیاری برای انجام نهان‌نگاری تصویر وجود ندارد در حالی که بیشتر آنها از CelebA، ImageNet یا BOSSBase استفاده می‌کنند. هر یک از روش‌ها روش‌ها و معیارهای ارزیابی خاص خود را دارند و از این رو بستر مشترکی برای مقایسه وجود ندارد. مقایسه ارزش پیک سیگنال به نویز (PSNR) که در جدول ۵ نشان داده شده است، ایده‌ای در مورد عملکرد امنیتی روش‌های مختلف ارائه می‌دهد. تا حد زیادی، بهترین مقدار PSNR 64.7 توسط کوپوسامی^۱ و همکاران با استفاده از چرخه GAN به دست آمده است. [۴۹] روش‌های مبتنی بر GAN عملکرد امنیتی و ظرفیت پنهان بهتری دارند. GAN به طور گسترده مورد

^۱ Kuppasamy

استفاده و ترجیح داده شده‌ترین معماری نسبت به رمزگذار خودکار و روش‌های آماری است. روش‌های سنتی از امنیت کمتری برخوردار هستند، زیرا فقط به تشخیص وجود پیام مخفی بستگی دارد. پیام مخفی را می‌توان به راحتی استخراج کرد زیرا جاسازی از یک روش آماری استفاده می‌کند.

به طور خلاصه، این مقاله در مورد تکنیک‌های مورد استفاده در زمان‌های اخیر برای نهان‌نگاری تصویر، روندهای فعلی را توضیح داده است. همراه با آن، جزئیات مربوط به مجموعه داده‌ها و معیارهای ارزیابی به تفصیل بیان شده است. چالش‌های پیش‌رو، بحث‌هایی درباره شکاف‌ها و زمینه‌های جهت‌گیری آینده نیز در این مقاله مورد ارزیابی قرار می‌گیرد. می‌توان نتیجه گرفت که یادگیری عمیق با توجه به پرشدن تمام چالش‌ها و شکاف‌ها، پتانسیل فوق العاده‌ای در زمینه نهان‌نگاری تصویر دارد.

منابع:

- [1] Wikipedia. (2020). Steganography. [Online]. Available: <https://en.wikipedia.org/wiki/Steganography>
- [2] H. Shi, X.-Y. Zhang, S. Wang, G. Fu, and J. Tang, "Synchronized detection and recovery of steganographic messages with adversarial learning," in Proc. Int. Conf. Comput. Sci. Cham, Switzerland: Springer, 2019, pp. 31-43.
- [3] N. F. Hordri, S. S. Yuhaniz, and S. M. Shamsuddin, "Deep learning and its applications: A review," in Proc. Conf. Postgraduate Annu. Res. Informat. Seminar, 2016, pp. 1-6.
- [4] N. F. Johnson and S. Jajodia, "Exploring steganography: Seeing the unseen," Computer, vol. 31, no. 2, pp. 26-34, Feb. 1998.
- [5] S. Gupta, G. Gujral, and N. Aggarwal, "Enhanced least significant bit algorithm for image steganography," Int. J. Comput. Eng. Manage., vol. 15, no. 4, pp. 40-42, 2012.
- [6] R. Das and T. Tuithung, "A novel steganography method for image based on Huffman encoding," in Proc. 3rd Nat. Conf. Emerg. Trends Appl. Comput. Sci., Mar. 2012, pp. 14-18.
- [7] A. Singh and H. Singh, "An improved LSB based image steganography technique for RGB images," in Proc. IEEE Int. Conf. Electr., Comput. Commun. Technol. (ICECCT), Mar. 2015, pp. 1-4.
- [8] Z. Qu, Z. Cheng, W. Liu, and X. Wang, "A novel quantum image steganography algorithm based on exploiting modification direction," Multimedia Tools Appl., vol. 78, no. 7, pp. 7981-8001, Apr. 2019.
- [9] S. Wang, J. Sang, X. Song, and X. Niu, "Least significant qubit (LSQb) information hiding algorithm for quantum image," Measurement, vol. 73, pp. 352-359, Sep. 2015.
- [10] N. Patel and S. Meena, "LSB based image steganography using dynamic key cryptography," in Proc. Int. Conf. Emerg. Trends Commun. Technol. (ETCT), Nov. 2016, pp. 1-5.
- [11] O. Elharrouss, N. Almaadeed, and S. Al-Maadeed, "An image steganography approach based on k-least significant bits (k-LSB)," in Proc. IEEE Int. Conf. Informat., IoT, Enabling Technol. (ICIOT), Feb. 2020, pp. 131-135.
- [12] M. V. S. Tarun, K. V. Rao, M. N. Mahesh, N. Srikanth, and M. Reddy, "Digital video steganography using LSB technique," Red, vol. 100111, Apr. 2020, Art. no. 11001001.
- [13] S. S. M. Than, "Secure data transmission in video format based on LSB and Huffman coding," Int. J. Image, Graph. Signal Process., vol. 12, no. 1, p. 10, 2020.
- [14] M. B. Tuieb, M. Z. Abdullah, and N. S. Abdul-Razaq, "An efficiency, secured and reversible video steganography approach based on least significant," J. Cellular Automata, vol. 16, no. 17, Apr. 2020.
- [15] R. J. Mstafa, K. M. Elleithy, and E. Abdelfattah, "A robust and secure video steganography method in DWT-DCT domains based on multiple object tracking and ECC," IEEE Access, vol. 5, pp. 5354-5365, 2017.
- [16] K. A. Al-Afandy, O. S. Faragallah, A. Elmhawly, E.-S.-M. El-Rabaie, and G. M. El-Banby, "High security data hiding using image cropping and LSB least significant bit steganography," in Proc. 4th IEEE Int. Colloq. Inf. Sci. Technol. (CiSt), Oct. 2016, pp. 400-404.
- [17] A. Arya and S. Soni, "Performance evaluation of secret image steganography techniques using least significant bit (LSB) method," Int. J. Comput. Sci. Trends Technol., vol. 6, no. 2, pp. 160-165, 2018.
- [18] G. Swain, "Very high-capacity image steganography technique using quotient value differencing and LSB substitution," Arabian J. Sci. Eng., vol. 44, no. 4, pp. 2995-3004, Apr. 2019.
- [19] A. Qiu, X. Chen, X. Sun, S. Wang, and W. Guo, "Coverless image steganography method based on feature selection," J. Inf. Hiding Privacy Protection, vol. 1, no. 2, p. 49, 2019.
- [20] R. D. Rashid and T. F. Majeed, "Edge based image steganography: Problems and solution," in Proc. Int. Conf. Commun., Signal Process., Appl. (ICCSPA), Mar. 2019, pp. 1-5.

- [21] X. Liao, J. Yin, S. Guo, X. Li, and A. K. Sangaiah, "Medical JPEG image steganography based on preserving inter-block dependencies," *Comput. Electr. Eng.*, vol. 67, pp. 320-329, Apr. 2018.
- [22] W. Lu, Y. Xue, Y. Yeung, H. Liu, J. Huang, and Y. Shi, "Secure halftone image steganography based on pixel density transition," *IEEE Trans. Dependable Secure Comput.*, early access, Aug. 6, 2019, doi: 10.1109/TDSC.2019.2933621.
- [23] Y. Zhang, C. Qin, W. Zhang, F. Liu, and X. Luo, "On the fault-tolerant performance for a class of robust image steganography," *Signal Process.*, vol. 146, pp. 99-111, May 2018.
- [24] H. M. Sidqi and M. S. Al-Ani, "Image steganography: Review study," in *Proc. Int. Conf. Image Process., Comput. Vis., Pattern Recognit. (IPCV)*, 2019, pp. 134-140.
- [25] P. Wu, Y. Yang, and X. Li, "Image-into-image steganography using deep convolutional network," in *Proc. Pacific Rim Conf. Multimedia*. Cham, Switzerland: Springer, 2018, pp. 792-802.
- [26] P. Wu, Y. Yang, and X. Li, "StegNet: Mega image steganography capacity with deep convolutional network," *Future Internet*, vol. 10, no. 6, p. 54, Jun. 2018.
- [27] X. Duan, K. Jia, B. Li, D. Guo, E. Zhang, and C. Qin, "Reversible image steganography scheme based on a U-Net structure," *IEEE Access*, vol. 7, pp. 9314-9323, 2019.
- [28] T. P. Van, T. H. Dinh, and T. M. Thanh, "Simultaneous convolutional neural network for highly efficient image steganography," in *Proc. 19th Int. Symp. Commun. Inf. Technol. (ISCIT)*, Sep. 2019, pp. 410-415.
- [29] R. Rahim and S. Nadeem, "End-to-end trained CNN encoder-decoder networks for image steganography," in *Proc. Eur. Conf. Comput. Vis. (ECCV)*, 2018, pp. 1-6.
- [30] Z. Wang, N. Gao, X. Wang, J. Xiang, and G. Liu, "STNet: A style transformation network for deep image steganography," in *Proc. Int. Conf. Neural Inf. Process*. Cham, Switzerland: Springer, 2019, pp. 3-14.
- [31] K. Yang, K. Chen, W. Zhang, and N. Yu, "Provably secure generative steganography based on autoregressive model," in *Proc. Int. Workshop Digit. Watermarking*. Cham, Switzerland: Springer, 2018, pp. 55-68.
- [32] S. Baluja, "Hiding images in plain sight: Deep steganography," in *Proc. Adv. Neural Inf. Process. Syst.*, 2017, pp. 2069-2079.
- [33] R. Zhang, S. Dong, and J. Liu, "Invisible steganography via generative adversarial networks," *Multimedia Tools Appl.*, vol. 78, no. 7, pp. 8559-8575, Apr. 2019.
- [34] S. Islam, A. Nigam, A. Mishra, and S. Kumar, "VStegNET: Video steganography network using spatio-temporal features and micro-bottleneck," in *Proc. BMVC*, Sep. 2019, p. 274.
- [35] I. Goodfellow, J. Pouget-Abadie, M. Mirza, B. Xu, D. Warde-Farley, S. Ozair, A. Courville, and Y. Bengio, "Generative adversarial nets," in *Proc. Adv. Neural Inf. Process. Syst.*, 2014, pp. 2672-2680.
- [36] D. Volkhonskiy, B. Borisenko, and E. Burnaev, "Generative adversarial networks for image steganography," in *Proc. ICRL Conf.*, 2016.
- [37] D. Volkhonskiy, I. Nazarov, and E. Burnaev, "Steganographic generative adversarial networks," in *Proc. 12th Int. Conf. Mach. Vis. (ICMV)*, vol. 11433, 2020, Art. no. 114333M.
- [38] D. J. Im, C. D. Kim, H. Jiang, and R. Memisevic, "Generating images with recurrent adversarial networks," 2016, arXiv:1602.05110. [Online]. Available: <http://arxiv.org/abs/1602.05110>
- [39] H. Shi, J. Dong, W. Wang, Y. Qian, and X. Zhang, "SSGAN: Secure steganography based on generative adversarial networks," in *Proc. Pacific Rim Conf. Multimedia*. Cham, Switzerland: Springer, 2017, pp. 534-544.
- [40] J. Hayes and G. Danezis, "Generating steganographic images via adversarial training," in *Proc. Adv. Neural Inf. Process. Syst.*, 2017, pp. 1954-1963.
- [41] J. Yang, K. Liu, X. Kang, E. K. Wong, and Y.-Q. Shi, "Spatial image steganography based on generative adversarial network," 2018, arXiv:1804.07939. [Online]. Available: <http://arxiv.org/abs/1804.07939>

- [42] J. Yang, D. Ruan, J. Huang, X. Kang, and Y.-Q. Shi, "An embedding cost learning framework using GAN," IEEE Trans. Inf. Forensics Security, vol. 15, pp. 839-851, 2020. 23422 VOLUME 9, 2021 N. Subramanian et al.: Image Steganography: A Review of the Recent Advances
- [43] W. Tang, S. Tan, B. Li, and J. Huang, "Automatic steganographic distortion learning using a generative adversarial network," IEEE Signal Process. Lett., vol. 24, no. 10, pp. 1547-1551, Oct. 2017.
- [44] H. Naito and Q. Zhao, "A new steganography method based on generative adversarial networks," in Proc. IEEE 10th Int. Conf. Awareness Sci. Technol. (iCAST), Oct. 2019, pp. 1-6.
- [45] K. Zhang, A. Cuesta-Infante, and K. Veeramachaneni, "SteganoGAN: Pushing the limits of image steganography," Jan. 2019, arXiv:1901.03892. [Online]. Available: <https://arxiv.org/abs/1901.03892>
- [46] J. Zhu, R. Kaplan, J. Johnson, and L. Fei-Fei, "Hidden: Hiding data with deep networks," in Proc. Eur. Conf. Comput. Vis. (ECCV), 2018, pp. 657-672.
- [47] Y. Ke, M. Zhang, J. Liu, T. Su, and X. Yang, "Generative steganography with Kerckhoffs' principle based on generative adversarial networks," 2017, arXiv:1711.04916. [Online]. Available: <http://arxiv.org/abs/1711.04916>
- [48] J.-Y. Zhu, T. Park, P. Isola, and A. A. Efros, "Unpaired image-to-image translation using cycle-consistent adversarial networks," in Proc. IEEE Int. Conf. Comput. Vis. (ICCV), Oct. 2017, pp. 2223-2232.
- [49] P. G. Kuppusamy, K. C. Ramya, S. Sheebha Rani, M. Sivaram, and V. Dhasarathan, "A novel approach based on modified cycle generative adversarial networks for image steganography," Scalable Comput., Pract. Exper., vol. 21, no. 1, pp. 6372, Mar. 2020.
- [50] C. Chu, A. Zhmoginov, and M. Sandler, "CycleGAN, a master of steganography," 2017, arXiv:1712.02950. [Online]. Available: <http://arxiv.org/abs/1712.02950>
- [51] H. Porav, V. Musat, and P. Newman, "Reducing steganography in cycle-consistency GANs," in Proc. CVPR Workshops, 2019, pp. 78-82.
- [52] R. Meng, Q. Cui, Z. Zhou, Z. Fu, and X. Sun, "A steganography algorithm based on CycleGAN for covert communication in the Internet of Things," IEEE Access, vol. 7, pp. 90574-90584, 2019.
- [53] A. Odena, C. Olah, and J. Shlens, "Conditional image synthesis with auxiliary classifier GANs," in Proc. Int. Conf. Mach. Learn., 2017, pp. 2642-2651.
- [54] Z. Zhang, G. Fu, J. Liu, and W. Fu, "Generative information hiding method based on adversarial networks," in Proc. Int. Conf. Comput. Eng. Netw. Cham, Switzerland: Springer, 2018, pp. 261-270.
- [55] M.-M. Liu, M.-Q. Zhang, J. Liu, Y.-N. Zhang, and Y. Ke, "Coverless information hiding based on generative adversarial networks," 2017, arXiv:1712.06951. [Online]. Available: <http://arxiv.org/abs/1712.06951>
- [56] X. Duan, H. Song, C. Qin, and M. K. Khan, "Coverless steganography for digital images based on a generative model," Comput., Mater. Continua, vol. 55, no. 3, pp. 483-493, Jul. 2018.
- [57] M. Arjovsky, S. Chintala, and L. Bottou, "Wasserstein GAN," 2017, arXiv:1701.07875. [Online]. Available: <http://arxiv.org/abs/1701.07875>
- [58] C. Li, Y. Jiang, and M. Cheslyar, "Embedding image through generated intermediate medium using deep convolutional generative adversarial network," Comput., Mater. Continua, vol. 56, no. 2, pp. 313-324, 2018.
- [59] Z. Wang, N. Gao, X. Wang, X. Qu, and L. Li, "SStEGAN: Self-learning steganography based on generative adversarial networks," in Proc. Int. Conf. Neural Inf. Process. Cham, Switzerland: Springer, 2018, pp. 253-264.
- [60] P. Bas, T. Filler, and T. Pevný, "'Break our steganographic system': The ins and outs of organizing BOSS," in Information Hiding, T. Filler, T. Pevný, S. Craver, and A. Ker, Eds. Berlin, Germany: Springer, 2011, pp. 59-70.
- [61] T. Pevný, T. Filler, and P. Bas, "Using high-dimensional image models to perform highly undetectable steganography," in Information Hiding, R. Böhme, P. W. L. Fong, and R. Safavi-Naini, Eds. Berlin, Germany: Springer, 2010, pp. 161-177.

- [62] Z. Liu, P. Luo, X. Wang, and X. Tang, "Deep learning face attributes in the wild," in Proc. IEEE Int. Conf. Comput. Vis. (ICCV), Dec. 2015, pp. 3730-3738.
- [63] J. Deng, W. Dong, R. Socher, L.-J. Li, K. Li, and L. Fei-Fei, "ImageNet: A large-scale hierarchical image database," in Proc. IEEE Conf. Comput. Vis. Pattern Recognit., Jun. 2009, pp. 248-255.
- [64] L. Deng, "The MNIST database of handwritten digit images for machine learning research [best of the Web]," IEEE Signal Process. Mag., vol. 29, no. 6, pp. 141-142, Nov. 2012.
- [65] T.-Y. Lin, M. Maire, S. Belongie, L. Bourdev, R. Girshick, J. Hays, P. Perona, D. Ramanan, C. L. Zitnick, and P. Dollár, "Microsoft COCO: Common objects in context," in Proc. Eur. Conf. Comput. Vis., 2014, pp. 740-755.
- [66] E. Agustsson and R. Timofte, "NTIRE 2017 challenge on single image super-resolution: Dataset and study," in Proc. IEEE Conf. Comput. Vis. Pattern Recognit. Workshops (CVPRW), Jul. 2017, pp. 126-135.
- [67] W. Tang, S. Tan, B. Li, and J. Huang, "Automatic steganographic distortion learning using a generative adversarial network," IEEE Signal Process. Lett., vol. 24, no. 10, pp. 1547-1551, Oct. 2017.
- [68] M. Zhang and A. A. Sawchuk, "USC-HAD: A daily activity dataset for ubiquitous activity recognition using wearable sensors," in Proc. ACM Int. Conf. Ubiquitous Comput. (UbiComp), Workshop Situation, Activity Goal Awareness (SAGAware), Pittsburgh, PA, USA, Sep. 2012, pp. 1036-1043.
- [69] M. Cimpoi, S. Maji, I. Kokkinos, S. Mohamed, and A. Vedaldi, "Describing textures in the wild," in Proc. IEEE Conf. Comput. Vis. Pattern Recognit. (CVPR), Jun. 2014, pp. 3606-3613.
- [70] B. G. Huang, M. Ramesh, T. Berg, and E. Learned-Miller, "Labeled faces in the wild: A database for studying face recognition in unconstrained environments," Univ. Massachusetts, Amherst, Amherst, MA, USA, Tech. Rep. 07-49, Oct. 2007.
- [71] M. Everingham, L. Van Gool, C. K. I. Williams, J. Winn, and A. Zisserman, "The Pascal visual object classes (VOC) challenge," Int. J. Comput. Vis., vol. 88, no. 2, pp. 303-338, Jun. 2010.
- [72] V. Holub and J. Fridrich, "Designing steganographic distortion using directional filters," in Proc. IEEE Int. Workshop Inf. Forensics Secur. (WIFS), Dec. 2012, pp. 234-239.