



ORIGINAL RESEARCH PAPER

A Novel Approach to Reducing Fraud in Insurance Processes Using GAN-Based Steganography and AES-GCM Encryption

Mona Parastesh^{1,*}¹ PhD Student, Department of Computer Engineering, Faculty of Technical and Engineering, Islamic Azad University, South Tehran Branch, Tehran, Iran.

ARTICLE INFO

Article History:

Received: 16 October 2025

Final revision: 20 January 2026

Accepted: 31 January 2026

Early online access: 31 January 2026

Published: 1 April 2026

Keywords:

Claim reporting

Deep learning

Information hiding

Insurance

Insurance fraud

ABSTRACT

BACKGROUND AND OBJECTIVES: One of the major challenges in the insurance industry is combating fraud during the claim submission process, which is particularly evident in digital claim systems where customers send images of damages and injuries through mobile applications. As the insurance sector continues its digital transformation, introducing modern methods across various operational stages such as underwriting, issuance, and claim processing, the necessity of developing innovative and secure fraud detection mechanisms has become more crucial than ever. With the rapid digitization and digitalization of insurance services, ensuring the authenticity of transmitted data and protecting sensitive information have emerged as key priorities. In digital claim reporting, where policyholders upload damage photographs via mobile applications, the authenticity of spatial and temporal metadata—specifically the location and time of the captured images—plays a vital role in reducing fraudulent submissions. To address this issue, this study proposes a novel security framework that utilizes steganography and cryptography for embedding encrypted metadata directly into the image file. In this method, the location and timestamp information extracted from the damage images are first encrypted using the AES-GCM algorithm and then securely embedded into the image through a Generative Adversarial Network (GAN)-based steganographic model. The hidden data can subsequently be extracted and decrypted at the insurer's end for verification. This approach not only strengthens the authenticity verification of the submitted images but also enhances data confidentiality and resistance to tampering. By integrating this method into a company's insurance application, insurers can build greater trust with policyholders, automate the verification process, and effectively reduce costs associated with fraudulent activities. The combination of deep learning-based steganography with robust encryption offers a promising path toward intelligent, secure, and fraud-resistant insurance claim systems.

METHODS: The proposed system employs a GAN-based steganographic framework to embed spatial and temporal metadata within damage images. Initially, the Exif metadata (including the capture time and GPS coordinates) is extracted and encrypted using the AES-GCM (Advanced Encryption Standard – Galois/Counter Mode) algorithm. This encryption stage ensures confidentiality and integrity of sensitive data before it is hidden in the image. The encrypted data are then embedded using a deep neural network generator, which learns to conceal information within the least perceptible image regions while preserving high visual fidelity. For performance evaluation, the system was tested on a dataset of 100 automobile damage photographs having been collected from real insurance claim records. Three methods were compared: Baseline LSB (Least Significant Bit) method, GAN-based steganography without encryption, and GAN combined with AES-GCM encryption (proposed method). The evaluation considered several quantitative metrics: Image quality measured by Peak Signal-to-Noise Ratio (PSNR) and Structural Similarity Index (SSIM), Embedding capacity, Bit Error Rate (BER) under standard steganalysis attacks such as JPEG compression (quality 50), Gaussian noise ($\sigma = 5$), and image scaling (0.9), and Detection rate of metadata manipulation to assess system robustness against tampering.

FINDINGS: The experimental results confirmed the effectiveness of the proposed approach. Images containing hidden data maintained a high visual quality with an average PSNR of 41 dB, demonstrating that the steganographic embedding had minimal perceptual impact. The GAN-based method outperformed the baseline LSB technique, achieving higher image quality (40.6 ± 0.8 dB vs. 36.7 ± 1.1 dB) and better structural similarity. Moreover, the average BER under compression and noise attacks was nearly half that of the baseline method ($\approx 12\text{--}19\%$ vs. $\approx 29\text{--}34\%$). Integrating AES-GCM encryption had no negative effect on visual quality but significantly improved end-to-end data security. Statistical analysis using a paired-sample t-test revealed that the performance differences were significant at a 95% confidence level ($p < 0.001$). The 95% confidence interval for PSNR values of the proposed model was calculated to be [40.7, 41.3], reflecting the system's stability and consistency. In the metadata tampering experiments, all instances of manipulated metadata in the 100-image dataset were successfully detected, achieving 100% detection accuracy. The average mobile implementation time for the complete encryption and embedding sequence was approximately 2.0 seconds per image, confirming the feasibility of applying the system in real-time insurance applications.

CONCLUSION: The findings of this research demonstrate that employing deep learning-based steganography (GAN) in combination with AES-GCM encryption can effectively mitigate security challenges in online insurance claim processes. The proposed method provides a reliable mechanism to verify the authenticity of spatial and temporal data while maintaining high image quality and computational efficiency. This integrated framework enhances both data protection and fraud prevention, ensuring that insurance companies can verify the legitimacy of submitted claim images with greater confidence. Consequently, the proposed GAN-AES hybrid model achieves an optimal balance between image fidelity, robustness to steganographic attacks, and metadata integrity verification. The results suggest that this approach can serve as a practical and efficient layer of protection for digital insurance claim systems, contributing to enhancing trust, transparency, and security across the insurance ecosystem.

*Corresponding Author:

Email: St_m_parastesh@azad.ac.ir

Phone: +98 2129465229

ORCID: 0009-0008-2851-7702

DOI: [10.22056/ijir.2026.02.02](https://doi.org/10.22056/ijir.2026.02.02)This is an open access article under the CC BY license (<http://creativecommons.org/licenses/by/4.0/>).



مقاله پژوهشی

رویکرد نوین برای کاهش تقلب در فرایندهای بیمه‌ای با استفاده از پنهان‌نگاری مبتنی بر شبکه‌های مولد متخاصم (GAN) و رمزنگاری AES-GCM

مونا پرستش^{۱*}

^۱ دانشجوی دکتری، گروه مهندسی کامپیوتر، دانشکده فنی و مهندسی، دانشگاه آزاد اسلامی، واحد تهران جنوب، تهران، ایران

چکیده:

پیشینه و اهداف: یکی از چالش‌های بزرگ در صنعت بیمه، مقابله با تقلب در فرایند اعلام خسارت‌هاست که به‌ویژه در روش‌های دیجیتال، از جمله ارسال عکس‌های خسارت از طریق اپلیکیشن‌ها، به چشم می‌خورد و به نظر می‌رسد صنعت بیمه بیش از هر زمان دیگری به بررسی روش‌های کشف تقلب نوین و نیز ایمن‌سازی اطلاعات نیاز دارد. همچنین امنیت و اصالت اطلاعات مکانی و زمانی تصاویر ارسالی از طریق اپلیکیشن نقش مهمی در کاهش تقلب دارد. در این پژوهش، روشی نوین ارائه شده است که با استفاده از تکنیک پنهان‌نگاری، اطلاعات مکانی و زمانی عکس‌های ارسالی از سوی بیمه‌گذاران رمزنگاری و مستقیماً در فایل تصویر جاسازی شده، سپس در مقصد رمزگشایی می‌شود. این رویکرد نه تنها امکان احراز اصالت اطلاعات ارائه‌شده را افزایش می‌دهد، بلکه امنیت داده‌ها را نیز تضمین می‌کند و خطر تقلب را کاهش می‌دهد. پیاده‌سازی این روش می‌تواند به بهبود اعتماد بین بیمه‌گر و بیمه‌گذار و کاهش هزینه‌های ناشی از تخلفات کمک کند.

روش‌شناسی: در این پژوهش، از رویکرد استگنوگرافی مبتنی بر شبکه‌های مولد متخاصم (GAN) برای جاسازی اطلاعات مکانی و زمانی تصاویر خسارت استفاده شده است ابتدا اطلاعات مربوط به زمان و مکان عکس با استفاده از الگوریتم رمزنگاری AES-GCM رمزنگاری شده تا امنیت داده‌های حساس تضمین شود. برای ارزیابی سیستم، داده‌های پنهان‌شده از تصاویر استخراج و با داده‌های اصلی مقایسه شدند. سیستم روی یک مجموعه داده ۱۰۰ تایی از تصاویر خسارت بیمه‌ای خودرو آزمایش و سه روش مقایسه‌ای بررسی شد. روش پایه (LSB)، استگنوگرافی مبتنی بر شبکه‌های مولد متخاصم، و ترکیب شبکه‌های مولد متخاصم با رمزنگاری AES-GCM. معیارهای ارزیابی شامل کیفیت تصویر (PSNR)، ظرفیت پنهان‌سازی (SSIM)، میانگین خطای بیت تحت حملات استاندارد و نرخ تشخیص دستکاری در فراداده‌ها محاسبه شد.

یافته‌ها: در این پژوهش مشخص شد تصاویر حاوی داده‌های پنهان میانگین PSNR برابر با ۴۱ dB داشتند که نشان‌دهنده حفظ کیفیت بصری بالای تصاویر است. روش شبکه‌های مولد متخاصم، میانگین $PSNR \approx 40.6 \pm 0.8$ dB (در مقابل ۳۶.۷ ± 1.1 dB برای روش پایه) و SSIM بالاتری نشان داد؛ میانگین BER در برابر فشرده‌سازی و نویز برای شبکه‌های مولد متخاصم تقریباً نصف روش پایه بود (به ترتیب ۱۲-۱۹٪ در برابر ۲۹-۳۴٪). افزودن AES-GCM تأثیری بر کیفیت تصویری نداشت، اما امنیت انتها-به-انتها را افزایش داد. آزمون t زوجی نشان داد این تفاوت از نظر آماری در سطح اطمینان ۹۵٪ معنادار است ($p < 0.001$). در آزمون‌های دستکاری فراداده در آزمایش حاضر، در مجموعه داده آزمایشی ۱۰۰ تصویری، تمامی موارد دستکاری فراداده (meta data) با موفقیت شناسایی شد. زمان پردازش میانگین پیاده‌سازی موبایلی برای کل زنجیره رمزنگاری-جاسازی تقریباً ۲ ثانیه در هر تصویر بود.

نتیجه‌گیری: این پژوهش نشان داد که استفاده از استگنوگرافی مبتنی بر یادگیری عمیق می‌تواند چالش‌های امنیتی در فرایند اعلام خسارت آنلاین را برطرف کند. روش پیشنهادی ضمن تضمین کیفیت و امنیت تصاویر، امکان تأیید صحت اطلاعات مکانی و زمانی را فراهم می‌کند و می‌تواند به کاهش تقلب در بیمه کمک شایان توجهی کند. نتایج نشان می‌دهد که ترکیب استگنوگرافی مبتنی بر روش پیشنهادی با AES-GCM بهترین تعادل بین کیفیت تصویری، پایداری در برابر حملات و قابلیت تشخیص دستکاری فراداده را فراهم می‌آورد.

اطلاعات مقاله

تاریخ‌های مقاله:

دریافت: ۲۴ مهر ۱۴۰۴
بازنگری نهایی: ۳۰ دی ۱۴۰۴
پذیرش: ۱۱ بهمن ۱۴۰۴
زودآیند: ۱۱ بهمن ۱۴۰۴
انتشار: ۱۲ فروردین ۱۴۰۵

کلمات کلیدی:

اعلام خسارت
بیمه
پنهان‌سازی اطلاعات
تقلب بیمه‌ای
یادگیری عمیق

*نویسنده مسئول:

ایمیل: St_m_parastesh@azad.ac.ir

تلفن: ۰۹۸ ۲۱۲۹۴۶۵۲۲۹

ORCID: 0009-0008-2851-7702

DOI: 10.22056/ijir.2026.02.02

توجه: مدت‌زمان بحث و انتقاد برای این مقاله تا ۱ جولای ۲۰۲۶ در وب‌سایت IJIR در «نمایش مقاله» باز است.

قابلیت را دارد که داده‌های رمزنگاری شده را بدون کاهش محسوس کیفیت بصری در تصویر تعبیه کند.

مبانی نظری پژوهشی

کشف تقلب

تقلب یکی از معضلات اساسی در بسیاری از حوزه‌های اقتصادی و اجتماعی است که می‌تواند باعث اختلال در نظام‌های مالی، کاهش اعتماد عمومی و تحمیل هزینه‌های سنگین به سازمان‌ها و نهادها شود. به همین دلیل، شناسایی و کشف تقلب همواره به‌عنوان یکی از اهداف اصلی در نظام‌های نظارتی، مالی و اطلاعاتی مطرح بوده است. در این راستا، توسعه روش‌های هوشمند برای کشف الگوهای رفتاری غیرعادی یا ناسازگار با داده‌های واقعی، به حوزه‌ای مهم در علوم داده و فناوری اطلاعات تبدیل شده است.

یکی از بخش‌هایی که بیش از دیگر حوزه‌ها در معرض خطر تقلب قرار دارد، صنعت بیمه است. در این صنعت، به‌ویژه در فرایند رسیدگی به خسارت‌ها، تقلب می‌تواند در قالب بزرگ‌نمایی خسارت، ارائه مدارک جعلی، یا ساختن رخدادهای ساختگی رخ دهد. از این‌رو، فرایند کشف تقلب در بیمه با هدف شناسایی و مقابله با این رفتارهای فریبکارانه طراحی و استفاده می‌شود.

فرایند کشف تقلب در خسارت بیمه‌ای

فرایند کشف تقلب در حوزه بیمه، به مجموعه‌ای از روش‌ها و ابزارها اطلاق می‌شود که برای شناسایی و جلوگیری از ادعاهای جعلی یا گمراه‌کننده توسط بیمه‌گذاران طراحی شده‌اند. این فرایند می‌تواند شامل مراحل زیر باشد:

- **جمع‌آوری داده‌ها:** اطلاعات مربوط به خسارت، موقعیت مکانی، مستندات تصویری، سابقه بیمه‌گذار و تاریخچه خسارت‌ها ثبت می‌شوند.

- **پیش‌پردازش و تحلیل داده‌ها:** داده‌ها پاک‌سازی و برای تحلیل آماده می‌شوند. در این مرحله از روش‌های داده‌کاوی و یادگیری ماشین استفاده می‌شود.

- **شناسایی الگوهای مشکوک:** الگوریتم‌های تحلیل رفتاری یا شبکه‌های عصبی می‌توانند الگوهای غیرعادی را تشخیص دهند، مانند ادعاهای متعدد در بازه زمانی کوتاه، یا تطابق نداشتن مکان ثبت‌شده با اطلاعات واقعی.

- **ارزیابی و تصمیم‌گیری انسانی:** موارد مشکوک به تحلیلگران ارجاع داده می‌شوند تا با بررسی دقیق‌تر، درباره تقلبی بودن آن‌ها تصمیم‌گیری شود.

- **اقدامات پیشگیرانه و اصلاحی:** در صورت تأیید تقلب، اقدامات قانونی یا بیمه‌ای اعمال و سیستم به‌روزرسانی می‌شود تا در آینده هوشمندتر عمل کند (Viaene & Dedene, 2004).

یکی از چالش‌های اصلی در این فرایند، دشواری در تأیید صحت مستندات مانند تصاویر خسارت و محل وقوع حادثه است. روش پیش‌رو، روشی نوین برای شناسایی و برخورد با تقلب در ارسال

استگانوگرافی، که از کلمه‌ای یونانی به معنای «نوشتن پنهان» گرفته شده است، یکی از قدیمی‌ترین و مؤثرترین روش‌های امنیت اطلاعات است. این تکنیک با استفاده از رسانه‌هایی مانند تصاویر، ویدئوها، صداها، یا حتی متون، امکان پنهان‌نگاری داده‌ها را فراهم می‌کند، به‌گونه‌ای که وجود داده‌های پنهان برای افراد غیرمجاز غیرقابل تشخیص باشد. استگانوگرافی به‌ویژه در دنیای امروز که انتقال داده‌ها از طریق بسترهای دیجیتال امری رایج است، به ابزاری اساسی برای حفظ محرمانگی و امنیت اطلاعات تبدیل شده است (Mandal et al., 2022).

با وجود پیشرفت‌های متعدد در زمینه رمزنگاری و امنیت سایبری، استگانوگرافی همچنان جایگاه منحصربه‌فردی در زمینه پنهان‌سازی اطلاعات دارد. دلیل آن نه تنها توانایی این تکنیک در پنهان‌سازی داده‌ها بدون جلب توجه، بلکه ترکیب آن با سایر فناوری‌های پیشرفته مانند یادگیری ماشین و شبکه‌های عصبی است. این تکنیک در زمینه‌های مختلفی از جمله امنیت نظامی، پزشکی، مخابرات، و حتی سیستم‌های تجاری مانند بانکداری و بیمه به کار گرفته شده است.

در صنعت بیمه، یکی از چالش‌های اصلی، تقلب در فرایندهای اعلام خسارت است. این تقلب‌ها معمولاً شامل ارائه اطلاعات نادرست یا تغییر داده‌های موجود است، که می‌تواند به خسارت مالی سنگین برای شرکت‌های بیمه منجر شود. در سالیان گذشته روش‌های گوناگونی برای مبارزه با خسارت در بیمه‌های مختلف به کار گرفته شده است (احمدلو و همکاران، ۱۴۰۰). در حال حاضر، رویکردهای نوینی برای شفاف‌سازی و افزایش صحت داده‌های ثبت‌شده لازم است. استگانوگرافی در این زمینه می‌تواند به‌مثابه راهکاری مؤثر عمل کند؛ مثلاً، با جاسازی اطلاعات زمان و مکان در تصاویر خسارت، امکان جعل یا دستکاری داده‌ها کاهش می‌یابد.

هدف اصلی این پژوهش، استفاده از استگانوگرافی و رمزنگاری برای مقابله با این چالش در صنعت بیمه است. با ظهور فناوری‌های یادگیری عمیق، به‌ویژه شبکه‌های مولد متخاصم قابلیت‌های استگانوگرافی به میزان قابل توجهی بهبود یافته است. شبکه‌های مولد متخاصم، که ابتدا برای تولید تصاویر مصنوعی با کیفیت بالا توسعه یافته‌اند، اکنون به‌عنوان ابزاری قدرتمند برای تولید داده‌های پنهان نیز استفاده می‌شوند. مثلاً، پژوهش‌خوان و همکاران در سال ۲۰۲۰ نشان داده است که استفاده از شبکه‌های مولد متخاصم در کشف تقلب می‌تواند به کاهش خطاهای شناسایی و افزایش امنیت فرایندهای دیجیتال منجر شود (Khan et al., 2020).

در این پژوهش، چهارچوب ترکیبی نوینی پیشنهاد شده است که در آن اطلاعات زمان و مکان تصاویر خسارت بیمه‌ای ابتدا با استفاده از الگوریتم رمزنگاری AES-GCM رمزگذاری و سپس توسط یک شبکه مولد متخاصم در تصویر جاسازی می‌شود. استفاده از AES-GCM علاوه بر محرمانگی، امکان تشخیص هرگونه دستکاری در فراداده را نیز فراهم می‌آورد. در مقابل، شبکه‌های مولد متخاصم این

و نوع داده پنهان شده بستگی دارد. در ادامه، هریک از این رسانه‌ها و روش‌های پنهان‌سازی داده در آن‌ها توضیح داده می‌شود.

در تصاویر دیجیتال^۱، اطلاعات پنهان از طریق تغییرات جزئی در مقادیر پیکسل‌ها ذخیره می‌شود، به‌طوری که تغییرات ایجادشده برای چشم انسان غیرقابل تشخیص است. یکی از شناخته‌شده‌ترین روش‌ها در این زمینه، روش بیت کم‌اهمیت است که در آن بیت‌های انتهایی هر پیکسل با بیت‌های داده پنهان جایگزین می‌شوند. این روش از سادگی پیاده‌سازی و ظرفیت بالای جاسازی برخوردار است، هرچند در برابر فشرده‌سازی یا نویز ممکن است آسیب‌پذیر باشد (Subramanian et al., 2021).

در سیگنال‌های صوتی، داده‌ها معمولاً در دامنه زمانی یا فرکانسی سیگنال پنهان می‌شوند. به دلیل حساسیت بالای گوش انسان، پنهان‌سازی در حوزه صوتی چالش‌برانگیزتر است، اما با استفاده از مدل‌های ادراکی صوتی می‌توان داده‌ها را به‌نحوی جاسازی کرد که کیفیت سیگنال اصلی حفظ شود. این روش معمولاً در سیستم‌های چندرسانه‌ای، ارتباطات ایمن و کاربردهای حفاظتی استفاده می‌شود (Tushara & Navas, 2016).

در ویدئو^۲، فرایند پنهان‌سازی در فریم‌های متوالی انجام می‌شود. هر ویدئو شامل صدها یا هزاران فریم تصویری است، بنابراین، این رسانه ظرفیت بسیار بالایی برای پنهان‌سازی داده‌ها فراهم می‌کند. داده‌ها می‌توانند در فریم‌های خاص یا در بخش‌های حرکتی ویدئو قرار گیرند تا از تشخیص انسان و الگوریتم‌های تحلیلی جلوگیری شود. این روش به‌ویژه در حوزه‌هایی مانند ارتباطات محرمانه، پنهان‌سازی اطلاعات و امنیت محتوای دیجیتال کاربرد دارد (Khan et al., 2020).

در انواع متنی^۳، داده‌ها از طریق تغییرات بسیار ظریف در ساختار جملات، انتخاب واژگان، فاصله‌ها، یا نشانه‌گذاری‌ها پنهان می‌شوند. این نوع استگانوگرافی به دلیل ظرفیت محدود متن، بیشتر برای پیام‌های کوتاه یا رمزگذاری نشانه‌ای استفاده می‌شود. با وجود محدودیت ظرفیت، مزیت اصلی آن در سادگی انتشار و دشواری تشخیص از سوی ابزارهای خودکار است.

روش‌های متداول استگانوگرافی

استگانوگرافی یا داده پنهانی، علمی است که با هدف مخفی‌سازی اطلاعات درون رسانه‌های دیجیتال به‌گونه‌ای به کار می‌رود که وجود اطلاعات برای مشاهده‌گر یا مهاجم غیرقابل تشخیص باشد. کارایی هر روش استگانوگرافی براساس سه معیار اصلی ارزیابی می‌شود:

- **ظرفیت:** میزان داده‌ای که می‌توان بدون تخریب محسوس رسانه میزبان پنهان کرد؛
- **پایداری:** مقاومت داده‌های پنهان در برابر پردازش‌های رایج، مانند فشرده‌سازی، نویز و تغییر اندازه؛
- **غیرقابل تشخیص بودن:** توانایی روش در پنهان‌سازی

عکس‌های خسارت آنلاین در بیمه است، که برای جلوگیری و کشف عکس‌های تقلبی، و حفظ امنیت بیشتر طراحی شده است. استفاده از تکنیک‌های مکمل مانند رمزنگاری و استگانوگرافی در ثبت و صحت‌سنجی داده‌ها می‌تواند ارزش افزوده‌ای در فرایند کشف تقلب ایجاد کند.

استگانوگرافی (پنهان‌نگاری)

استگانوگرافی از واژه یونانی «Steganos» به معنای «پنهان» و «Graphy» به معنای «نوشتن» گرفته شده است. به‌طور کلی، استگانوگرافی علمی است که به پنهان‌سازی اطلاعات درون یک رسانه پوششی (مانند تصویر، صدا، ویدئو یا متن) می‌پردازد، به‌گونه‌ای که وجود اطلاعات پنهان‌شده برای دیگران غیرقابل شناسایی باشد. برخلاف رمزنگاری که داده‌ها را به شکلی ناخوانا تغییر می‌دهد، استگانوگرافی داده‌ها را به‌گونه‌ای در رسانه جاسازی می‌کند که حضور آن‌ها پنهان باقی بماند (Kadhim et al., 2019).

در نتیجه، پنهان‌نگاری به فن مخفی کردن اطلاعات درون فایل‌های دیجیتال گفته می‌شود، به‌طوری که تغییرات به‌طور غیرقابل مشاهده باقی بمانند. یکی از رایج‌ترین تکنیک‌ها در این زمینه، استفاده از تغییرات در بیت‌های کم‌اهمیت (LSB) است. در این روش، تغییرات جزئی در مقادیر باینری پیکسل‌های تصویر ایجاد می‌شود که برای چشم انسان قابل مشاهده نیستند. این تکنیک نه تنها در تصاویر، بلکه در فایل‌های صوتی و ویدئویی نیز استفاده می‌شود (Subramanian et al., 2021).

ساختار و عملکرد استگانوگرافی

فرایند استگانوگرافی به‌طور کلی از دو مرحله اصلی تشکیل می‌شود: **جاسازی داده‌ها و استخراج داده‌ها**. در مرحله جاسازی، داده‌های حساس درون یک رسانه پوششی (مانند تصویر دیجیتال) به‌گونه‌ای قرار می‌گیرند که تغییرات ایجادشده در ظاهر رسانه برای بیننده قابل تشخیص نباشد. ورودی این مرحله شامل رسانه پوششی (مانند تصویر اصلی) و داده‌های پنهان (مانند متن یا کلید رمزنگاری) است و خروجی آن، رسانه استگانوگرافیک است که داده‌های پنهان در آن جاسازی شده‌اند.

در مرحله استخراج، داده‌های پنهان شده از رسانه استگانوگرافیک بازیابی می‌شوند. در این فرایند، ورودی شامل رسانه حاوی داده پنهان و کلید مورد نیاز برای بازیابی است و خروجی، داده‌های اصلی استخراج‌شده خواهد بود (Rahman et al., 2020).

انواع رسانه‌های پوششی در استگانوگرافی

استگانوگرافی را می‌توان در طیف گسترده‌ای از رسانه‌ها به کار گرفت که هریک ویژگی‌ها، مزایا و چالش‌های خاص خود را دارند. رسانه‌های مورد استفاده معمولاً شامل **تصاویر دیجیتال**، **سیگنال‌های صوتی**، **ویدئوها** و **متون نوشتاری** هستند. انتخاب نوع رسانه به هدف کاربرد، میزان امنیت مورد نیاز، ظرفیت جاسازی

1 Image Steganography

2 Video Steganography

3 Text Steganography

داده‌ها، به‌طوری که تغییرات از دید انسان یا الگوریتم‌های آماری قابل تشخیص نباشند (سخنور و بابایی، ۱۳۹۵).
براساس نوع فضای جاسازی و الگوریتم‌های مورد استفاده، روش‌های استگانوگرافی را می‌توان در سه گروه اصلی و یک گروه ترکیبی طبقه‌بندی کرد:

۱. روش‌های مبتنی بر بیت کم/اهمیت (LSB)^۱

یکی از ساده‌ترین و پرکاربردترین روش‌ها در استگانوگرافی دیجیتال، روش LSB است. در این روش، داده‌ها در بیت‌های کم‌اهمیت پیکسل‌ها یا نمونه‌های صوتی جاسازی می‌شوند. از آنجاکه تغییر در بیت کم‌اهمیت باعث تغییرات بسیار جزئی در شدت رنگ یا صدا می‌شود، کیفیت ظاهری رسانه تقریباً ثابت می‌ماند و این روش از نظر سادگی، سرعت و ظرفیت بالا بسیار مورد توجه است.
بالین‌حال، روش LSB معایبی نیز دارد:

در برابر فشرده‌سازی، افزودن نویز یا تغییر اندازه تصویر بسیار آسیب‌پذیر است و به دلیل ماهیت مستقیم و خطی خود، در برابر تحلیل‌های آماری و حملات استگانوآنالیز ضعیف عمل می‌کند.
برای افزایش امنیت در این روش، محققان از الگوهای تصادفی برای انتخاب پیکسل‌ها یا رمزگذاری پیش از جاسازی استفاده می‌کنند تا الگوی پنهان‌سازی قابل شناسایی نباشد (آبرومند و همکاران، ۱۳۹۴).

۲. روش‌های مبتنی بر تبدیل^۲

برخلاف روش‌های مکانی، در روش‌های مبتنی بر تبدیل، داده‌ها در دامنه فرکانسی یا دامنه تبدیل‌شده تصویر جاسازی می‌شوند. متداول‌ترین تبدیل‌ها عبارت‌اند از تبدیل کسینوسی گسسته (DCT)، تبدیل موجک گسسته (DWT) و تبدیل فوریه گسسته (DFT).
این روش‌ها از پایداری بالاتری نسبت به LSB برخوردارند و در برابر فشرده‌سازی، فیلترگذاری و سایر حملات مقاوم‌تر هستند (Tushara & Navas, 2016).
بالین‌حال، این روش‌ها پیچیدگی محاسباتی بالاتری دارند و ممکن است ظرفیت جاسازی کمتری نسبت به روش‌های مکانی ارائه دهند.

۳. روش‌های مبتنی بر یادگیری عمیق^۳

در سال‌های اخیر، با ظهور هوش مصنوعی و یادگیری عمیق، نسل جدیدی از روش‌های استگانوگرافی معرفی شده است. این روش‌ها از شبکه‌های عصبی عمیق^۴ و شبکه‌های مولد متخاصم، برای یادگیری خودکار الگوی جاسازی استفاده می‌کنند.
در یک ساختار معمول مبتنی بر شبکه‌های مولد متخاصم، دو شبکه در تعامل رقابتی قرار دارند:

- شبکه مولد^۴، که هدف آن جاسازی داده‌ها در تصویر میزبان

است، به‌گونه‌ای که تشخیص آن غیرممکن باشد؛
• شبکه متمایزکننده^۵، که تلاش می‌کند تفاوت بین تصویر اصلی و تصویر جاسازی‌شده را تشخیص دهد.
در این فرایند، شبکه مولد به‌تدریج یاد می‌گیرد چگونه داده‌ها را به‌نحوی پنهان کند که حتی شبکه متمایزکننده نیز قادر به شناسایی آن نباشد. نتیجه این آموزش متقابل، تصاویری طبیعی و غیرقابل تشخیص از نمونه اصلی است که درعین‌حال حاوی داده‌های مخفی‌اند.

مزایای این روش‌ها عبارت‌اند از:
• ظرفیت بسیار بالا در جاسازی داده‌ها بدون افت کیفیت محسوس؛
• امنیت بالا در برابر تحلیل‌های آماری؛
• مقاومت در برابر فشرده‌سازی و نویز (Khan et al., 2020).
ازسوی دیگر، محدودیت‌هایی مانند نیاز به داده آموزشی زیاد، پیچیدگی محاسباتی و زمان آموزش طولانی نیز وجود دارد.

۴. روش‌های ترکیبی

به‌منظور ترکیب مزایای روش‌های یادشده و کاهش نقاط ضعف آن‌ها، پژوهشگران از روش‌های ترکیبی استفاده می‌کنند. مثلاً:
• ترکیب روش LSB با DWT یا DCT برای افزایش هم‌زمان ظرفیت و پایداری؛
• استفاده از GAN برای بهینه‌سازی محل جاسازی در حوزه تبدیل؛

• ترکیب رمزنگاری مانند (AES) با استگانوگرافی برای افزایش امنیت دولایه‌ای.
روش‌های ترکیبی معمولاً با هدف ایجاد تعادل بین ظرفیت، امنیت و پایداری طراحی می‌شوند و در کاربردهای حساس مانند امنیت اطلاعات پزشکی، تبادل داده‌های نظامی و کشف تقلب در بیمه بسیار پرکاربردند.

ویژگی‌های اصلی استگانوگرافی

استگانوگرافی به‌عنوان یکی از شاخه‌های مهم امنیت اطلاعات، دارای مجموعه‌ای از ویژگی‌های اصلی است که تعیین‌کننده کارایی و قابلیت اطمینان آن در کاربردهای مختلف‌اند. این ویژگی‌ها عبارت‌اند از:

۱. **محرمانگی:** اصلی‌ترین هدف استگانوگرافی، پنهان‌سازی وجود داده‌های حساس درون یک رسانه پوششی است. در یک سیستم استگانوگرافی مؤثر، حضور اطلاعات مخفی نباید توسط مشاهده‌گر انسانی یا ابزارهای آماری قابل تشخیص باشد. در واقع، وجود داده پنهان خود باید مخفی بماند.

۲. **امنیت:** در صورت شناسایی احتمال وجود داده پنهان، محتوای آن نباید قابل استخراج یا تفسیر باشد. این هدف معمولاً با ترکیب رمزنگاری و جاسازی مخفیانه محقق می‌شود. در این

1 Least Significant Bit

2 Transform-Based

3 DNN

4 Generator

5 Discriminator

در این نوع ارتباط، پیام‌ها در قالب داده‌های چندرسانه‌ای عادی (مانند تصاویر روزمره یا فایل‌های صوتی) پنهان شده و بدون جلب توجه منتقل می‌شوند. چنین قابلیتی در شرایطی که نیاز به ارتباطات امن و نامحسوس وجود دارد، اهمیت حیاتی دارد.

علاوه بر موارد بالا، استگنوگرافی در **حوزه‌های صنعتی و مالی** نیز به کار گرفته می‌شود. یکی از نمونه‌های نوین استفاده از این فناوری، **جلوگیری از تقلب در فرایندهای بیمه** است. در این کاربرد، اطلاعاتی مانند زمان، مکان و شناسه ثبت خسارت در تصاویر ارسالی توسط بیمه‌گذاران جاسازی می‌شود تا از ارسال تصاویر جعلی یا ویرایش‌شده جلوگیری شود. چنین رویکردی موجب افزایش اعتماد، کاهش خسارات ناشی از تقلب و ارتقای شفافیت در صنعت بیمه می‌شود.

به‌طور کلی، استگنوگرافی در حوزه‌هایی همچون امنیت سایبری، تبادل داده‌های حساس، کنترل اصالت اسناد دیجیتال و جلوگیری از جعل و تقلب، نقش اساسی ایفا می‌کند. این روش به‌ویژه در کاربردهایی که نیاز به حفظ محرمانگی اطلاعات هنگام انتقال از طریق اینترنت یا رسانه‌های عمومی وجود دارد، بسیار مؤثر و غیرقابل جایگزین است (Mandal et al., 2022).

شبکه‌های مولد متخاصم (SNAG)¹

شبکه‌های مولد متخاصم از پیشرفته‌ترین مدل‌های یادگیری عمیق هستند که نخستین بار در سال ۲۰۱۴ «یان گودفلو» و همکارانش آن را معرفی کردند (Goodfellow et al., 2014). هدف اصلی این شبکه‌ها، تولید داده‌های مصنوعی است که از داده‌های واقعی غیرقابل تشخیص باشند. شبکه‌های مولد متخاصم شامل دو شبکه عصبی متقابل با عملکرد رقابتی اند: **مولد**^۲ و **تفکیک‌کننده**^۳.
۱. **شبکه مولد**: این قسمت وظیفه دارد داده‌هایی تولید کند که شباهت زیادی به داده‌های واقعی داشته باشند. ورودی این شبکه معمولاً یک بردار نویز تصادفی است که از توزیع گاوسی یا یکنواخت نمونه‌برداری می‌شود. شبکه مولد طی فرایند آموزش، می‌کوشد الگوهای نهفته در داده‌های واقعی را بیاموزد و داده‌های جدیدی ایجاد کند که از دید انسان و حتی از دید الگوریتم‌های آماری، واقعی به نظر برسند.

۲. **شبکه تفکیک‌کننده** به‌عنوان یک طبقه‌بندی‌کننده عمل می‌کند و هدف آن تمایز میان داده‌های واقعی و داده‌های تولیدی است. این شبکه احتمال واقعی بودن ورودی‌ها را تخمین می‌زند و در واقع «قاضی» میان دو شبکه محسوب می‌شود. در طول فرایند آموزش، هر دو شبکه به‌صورت هم‌زمان و رقابتی به‌روزرسانی می‌شوند: مولد تلاش می‌کند داده‌هایی تولید کند که تفکیک‌کننده را فریب دهد، درحالی‌که تفکیک‌کننده سعی دارد توانایی تشخیص خود را افزایش دهد. این رقابت دوطرفه سبب می‌شود که مولد به مرور داده‌هایی با کیفیت و شباهت بصری بسیار بالا تولید کند، تا

حالت، حتی اگر مهاجم موفق به کشف داده پنهان شود، بدون کلید رمزنگاری شده قادر به درک محتوای آن نخواهد بود.

۳. **ظرفیت**: ظرفیت به میزان داده‌ای اشاره دارد که می‌توان در رسانه پوششی ذخیره کرد، بدون آنکه کیفیت یا ویژگی‌های آماری آن تغییر محسوسی پیدا کند. افزایش ظرفیت معمولاً به کاهش غیرقابل تشخیص بودن منجر می‌شود، بنابراین، انتخاب روش جاسازی باید بین این دو ویژگی تعادل برقرار کند.

۴. **مقاومت یا پایداری**: پایداری به توانایی حفظ داده‌های پنهان در برابر تغییرات یا حملات احتمالی روی رسانه مربوط است؛ از جمله فشرده‌سازی، اضافه شدن نویز، برش یا تغییر اندازه تصویر. یک روش مقاوم باید بتواند پس از این تغییرات نیز داده پنهان را به‌طور کامل بازیابی کند (Singh et al., 2016).

کاربردهای استگنوگرافی

استگنوگرافی در سال‌های اخیر به یکی از فناوری‌های کلیدی در حوزه امنیت داده و حفاظت از اطلاعات تبدیل شده است. این فناوری با فراهم کردن امکان پنهان‌سازی داده‌ها در انواع رسانه‌های دیجیتال، نقش مهمی در ایجاد ارتباطات امن و جلوگیری از دسترسی غیرمجاز ایفا می‌کند. از مهم‌ترین کاربردهای آن می‌توان به امنیت اطلاعات، احراز هویت دیجیتال، مبارزه با جعل، و ارتباطات مخفی اشاره کرد.

۱. **امنیت اطلاعات**: استگنوگرافی به‌عنوان روشی مکمل در کنار رمزنگاری به کار می‌رود تا لایه دوم امنیت داده‌ها را فراهم کند. در این روش، حتی اگر مهاجم به داده‌های رمزگذاری شده دست یابد، به‌دلیل پنهان بودن آن‌ها در یک رسانه پوششی (مانند تصویر یا صوت)، عملاً قادر به تشخیص یا بازیابی داده اصلی نخواهد بود. این ویژگی، استگنوگرافی را به ابزاری مؤثر در حفاظت از اطلاعات حساس در محیط‌های ناامن یا ارتباطات عمومی تبدیل کرده است (Mandal et al., 2022).

۲. **احراز هویت دیجیتال**: داده‌های تأییدکننده هویت مانند شماره سریال، شناسه کاربر یا امضای دیجیتال در تصاویر، ویدئوها یا اسناد جاسازی می‌شوند. این داده‌ها می‌توانند به‌صورت پنهان در فایل ذخیره شوند و فقط توسط سیستم‌های مجاز قابل استخراج باشند. به این ترتیب، اصالت هویت کاربران و منابع اطلاعاتی به‌صورت غیرقابل جعل حفظ می‌شود.

۳. **مبارزه با جعل و نقض مالکیت معنوی**: استگنوگرافی ابزاری کارآمد برای جاسازی اطلاعات مربوط به منشأ تولید محتوا، شناسه تولیدکننده یا کدهای ردیابی درون رسانه‌هاست. با استفاده از این روش، می‌توان منبع اصلی تولید داده را حتی پس از ویرایش یا انتشار مجدد اثر شناسایی کرد. این کاربرد به‌ویژه در حوزه‌هایی همچون حفاظت از آثار هنری دیجیتال، اسناد حقوقی و مدارک بیمه‌ای اهمیت ویژه‌ای دارد (Dalal & Juneja., 2021).

۴. **ارتباطات مخفی**: استگنوگرافی به‌عنوان ابزاری برای تبادل اطلاعات در محیط‌های نظارتی یا نظامی مورد استفاده قرار می‌گیرد.

1 Generative Adversarial Networks

2 Generator

3 Discriminator

بهتری دارند (Subramanian et al., 2021).

الگوریتم رمزنگاری

الگوریتم‌های رمزنگاری ابزارهای اساسی در حفظ امنیت داده‌ها هستند و نقش مهمی در محافظت از اطلاعات محرمانه، جلوگیری از نفوذ، و تضمین صحت ارتباطات دیجیتال ایفا می‌کنند. الگوریتم‌های رمزنگاری مجموعه‌ای از تکنیک‌های ریاضی هستند که برای حفاظت از محرمانگی، یکپارچگی، و صحت داده‌ها استفاده می‌شوند. این الگوریتم‌ها نقش مهمی در امنیت اطلاعات ایفا می‌کنند و به دو دسته اصلی رمزنگاری متقارن^۱ و رمزنگاری نامتقارن^۲ تقسیم می‌شوند.

رمزنگاری متقارن

در این روش، یک کلید واحد برای رمزنگاری و رمزگشایی استفاده می‌شود. این روش به دلیل سرعت بالا برای داده‌های بزرگ مناسب است، اما مشکل اصلی آن توزیع امن کلید میان طرفین است. الگوریتم AES^3 یکی از پرکاربردترین الگوریتم‌های متقارن است که امنیت بالایی را ارائه می‌دهد و برای داده‌های حجیم مناسب است. این الگوریتم در بسیاری از حوزه‌ها از جمله استeganوگرافی استفاده می‌شود تا امنیت داده‌های پنهان‌شده را تضمین کند (Kadhim et al., 2019).

رمزنگاری نامتقارن

در این روش، دو کلید جداگانه برای رمزنگاری و رمزگشایی استفاده می‌شود: یک کلید عمومی و یک کلید خصوصی. کلید عمومی برای رمزنگاری داده‌ها استفاده می‌شود، درحالی‌که تنها با کلید خصوصی می‌توان داده‌ها را رمزگشایی کرد. این روش معمولاً برای تبادل امن کلیدها یا داده‌های کوچک استفاده می‌شود. الگوریتم RSA^4 یکی از معروف‌ترین الگوریتم‌های رمزنگاری نامتقارن است که برای تضمین امنیت ارتباطات در شبکه‌ها استفاده می‌شود (Subramanian et al., 2021).

ویژگی‌های الگوریتم‌های رمزنگاری

الگوریتم‌های رمزنگاری بر پایه اصول ریاضی و منطق محاسباتی طراحی شده‌اند تا از دسترسی غیرمجاز، تغییر یا جعل داده‌ها جلوگیری کنند. مهم‌ترین ویژگی‌های کلیدی رمزنگاری عبارتند از:

- محرمانگی:** محرمانگی به این معناست که داده‌ها تنها برای افراد مجاز قابل دسترسی باشند و اشخاص یا سامانه‌های غیرمجاز نتوانند به محتوای اصلی اطلاعات دست پیدا کنند. این ویژگی از طریق استفاده از کلیدهای رمزنگاری و الگوریتم‌های ایمن (مانند AES در رمزنگاری متقارن یا RSA در رمزنگاری نامتقارن) محقق می‌شود.

جایی که حتی تفکیک‌کننده آموزش‌دیده نیز قادر به تمایز آن‌ها از داده‌های واقعی نباشد.

از دیدگاه یادگیری ماشین، این رقابت در چهارچوب یک بازی دوسویه کمینه-بیشینه مدل‌سازی می‌شود، که در آن تابع هزینه مولد و تفکیک‌کننده در جهت مخالف یکدیگر بهینه‌سازی می‌شوند. در نتیجه، زمانی که هر دو شبکه به تعادل برسند، داده‌های تولیدی مولد به قدری طبیعی خواهند بود که از داده‌های واقعی قابل تمایز نیستند.

از مهم‌ترین ویژگی‌های شبکه‌های مولد متخاصم می‌توان به آموزش متقابل، توانایی تولید داده‌های واقعی‌نما و انعطاف‌پذیری بالا در حوزه‌های مختلف اشاره کرد. شبکه‌های مولد متخاصم قادرند داده‌های پیچیده‌ای مانند تصاویر، ویدئوها، صداها و حتی متن‌های طبیعی را بازتولید کنند. در بسیاری از پژوهش‌ها، از این شبکه‌ها در زمینه‌هایی همچون تولید تصاویر مصنوعی، بازسازی تصاویر آسیب‌دیده، بهبود وضوح تصاویر، شبیه‌سازی داده‌های آموزشی، و همچنین در استگانوگرافی و امنیت اطلاعات استفاده شده است (Subramanian et al., 2021).

در حوزه استگانوگرافی مبتنی بر یادگیری عمیق، شبکه‌های مولد متخاصم کاربردی ویژه یافته‌اند. در این کاربرد، شبکه مولد به گونه‌ای آموزش داده می‌شود که داده‌های پنهان را درون تصویر میزبان جاسازی کند، به طوری که شبکه تفکیک‌کننده (یا هر الگوریتم تشخیص استگانوگرافی دیگر) نتواند تفاوتی میان تصویر اصلی و تصویر حاوی داده پنهان تشخیص دهد. این رویکرد سبب افزایش امنیت و ظرفیت پنهان‌سازی می‌شود و امکان جاسازی حجم بالایی از داده را بدون افت کیفیت بصری فراهم می‌آورد (Khan et al., 2020).

به طور مشابه، در مرحله استخراج داده‌ها، شبکه تفکیک‌کننده می‌تواند برای بازگردانی داده‌های پنهان‌شده از تصویر آموزش داده شود. در نتیجه، فرایند جاسازی و استخراج داده‌ها در چهارچوب یک سامانه یادگیری دوطرفه و هوشمند انجام می‌گیرد که دقت، سرعت و امنیت بالاتری نسبت به روش‌های سنتی مانند LSB ارائه می‌دهد.

کاربرد شبکه‌های مولد متخاصم در استگانوگرافی

شبکه‌های مولد متخاصم با فراهم‌سازی چارچوبی هوشمند و خودیادگیر، تحولی بنیادین در حوزه استگانوگرافی ایجاد کرده‌اند. این شبکه‌ها با ترکیب قابلیت‌های تولید داده واقعی‌نما و پنهان‌سازی هوشمند اطلاعات، مسیر توسعه نسل جدیدی از سامانه‌های امنیت داده و مقابله با تقلب را هموار کرده‌اند. از جمله مزایای استفاده از شبکه‌های مولد متخاصم در استگانوگرافی می‌توان به بهبود کیفیت بصری تصاویر استگانوگرافیک، افزایش مقاومت در برابر حملات تحلیل آماری و فشرده‌سازی، و افزایش ظرفیت جاسازی داده‌ها اشاره کرد. مطالعات متعددی نشان داده‌اند که روش‌های مبتنی بر شبکه‌های مولد متخاصم، در مقایسه با روش‌های کلاسیک مانند LSB یا تبدیل موجک، از نظر امنیت، پایداری و تشخیص‌ناپذیری عملکرد

1 Symmetric Cryptography

2 Asymmetric Cryptography

3 Advanced Encryption Standard

4 Rivest-Shamir-Adleman

کاراتسوباً^۳ همچنین تکنیک‌های مقاوم در برابر حملات کانال جانبی استفاده کنیم (Kim & Seo., 2025).

عملکرد و ویژگی‌های الگوریتم MCG-SEA
الگوریتم AES-GCM یکی از الگوریتم‌های رمزنگاری متقارن پیشرفته است که هم‌زمان رمزنگاری داده‌ها و احراز صحت آن‌ها را فراهم می‌کند. در این الگوریتم، داده‌های ورودی شامل متن اصلی، کلید متقارن معمولاً ۱۲۸ یا ۲۵۶ بیت، بردار مقدار اولیه و داده‌های اضافی قابل احراز هستند. داده‌های قابل اجرا شامل اطلاعاتی می‌شوند که رمزگذاری نمی‌شوند، اما صحت آن‌ها باید تأیید شود، مانند هدر پیام یا شناسه‌ها.

رمزگذاری در AES-GCM براساس حالت شمارنده انجام می‌شود؛ به این صورت که متن اصلی با استفاده از کلید و بردار اولیه رمزگذاری شده و خروجی آن داده‌های رمزنگاری شده هستند. استفاده از حالت شمارنده امکان رمزنگاری موازی و افزایش سرعت نسبت به حالت‌های سنتی مانند CBC^۴ را فراهم می‌کند. هم‌زمان با رمزنگاری، الگوریتم تگ احراز هویت را تولید می‌کند که معمولاً ۱۶ بیت طول دارد و با هدف تضمین یکپارچگی داده‌ها محاسبه می‌شود. تگ احراز هویت براساس متنی که رمزنگاری شده، بردار مقدار اولیه و داده‌های اضافی که رمز نمی‌شوند، تولید می‌شود. این تگ با استفاده از عملیات ضرب در میدان گالوا محاسبه می‌شود و تضمین می‌کند که هرگونه تغییر در داده‌های رمزنگاری شده یا کلید باعث عدم تطابق تگ خواهد شد و پیام به‌عنوان نامعتبر شناسایی می‌شود.

در فرایند رمزگشایی، گیرنده با استفاده از کلید، بردار اولیه و داده‌های رمزنگاری شده، مجدداً تگ را محاسبه می‌کند و در صورت تطابق با تگ دریافتی، پیام معتبر شناخته می‌شود؛ در غیر این صورت، پیام رد می‌شود (Kim & Seo., 2025).

AES-GCM علاوه‌بر فراهم آوردن سرعت بالا، امنیت و یکپارچگی داده‌ها را تضمین می‌کند و امکان شناسایی تغییرات، تکرار یا دستکاری پیام در طول انتقال را نیز فراهم می‌آورد.

به‌دلیل این ویژگی‌ها، AES-GCM به‌طور گسترده در سیستم‌های رمزنگاری مدرن، انتقال امن داده‌ها، ارتباطات شبکه‌ای و کاربردهای حساس به امنیت، مانند ذخیره‌سازی اطلاعات محرمانه و حفاظت از حریم خصوصی کاربران، استفاده می‌شود. این الگوریتم به‌دلیل ترکیب رمزگذاری سریع و تولید تگ احراز هویت، یکی از استانداردهای معتبر و مورد اعتماد در امنیت سایبری محسوب می‌شود.

در روش‌های ترکیبی، اعمال رمزنگاری پیش از انجام استگانوگرافی باعث افزایش مقاومت سیستم در برابر حملات مختلف استگانوگرافی می‌شود و امنیت داده‌های پنهان شده را به‌طور چشمگیری بهبود می‌بخشد (Singh et al., 2016). به‌ویژه در سامانه‌های بیمه، جایی که اطلاعات حساس شامل پرونده‌های خسارت، مستندات پزشکی، تصاویر و داده‌های شخصی بیمه‌گذاران مبادله می‌شوند، استفاده از

۲. یکپارچگی: این ویژگی تضمین می‌کند که داده‌ها در طول ذخیره‌سازی یا انتقال دچار تغییر، حذف یا دستکاری ناخواسته نشده باشند. برای تحقق این هدف، از روش‌هایی چون تابع درهم‌سازی و امضای دیجیتال استفاده می‌شود. در صورت تغییر هر بخش از پیام، هش تولیدی متفاوت خواهد بود و سیستم می‌تواند بلافاصله تغییر را تشخیص دهد.

۳. احراز هویت: احراز هویت به فرایندی اشاره دارد که در آن، هویت طرفین ارتباط (فرستنده و گیرنده) تأیید می‌شود. این ویژگی مانع از جعل هویت کاربران یا سامانه‌ها در فرایند تبادل داده می‌گردد. در پروتکل‌های رمزنگاری، احراز هویت از طریق گواهی‌های دیجیتال یا کلیدهای عمومی و خصوصی انجام می‌شود.

۴. غیرقابل انکار بودن: غیرقابل انکار بودن به این معناست که هیچ‌یک از طرفین نمی‌تواند مشارکت خود در ارسال یا دریافت داده‌ها را انکار کند. مثلاً، در تراکنش‌های مالی یا ارتباطات رسمی، امضای دیجیتال تضمین می‌کند که فرستنده پس از ارسال پیام، قادر به انکار آن نخواهد بود. این ویژگی نقش حیاتی در جلوگیری از اختلافات حقوقی و جعل اسناد الکترونیکی دارد (Tushara & Navas, 2016). این چهار ویژگی اساسی، پایه و اساس امنیت در تمامی سیستم‌های رمزنگاری مدرن به شمار می‌روند و در هر دو نوع رمزنگاری متقارن و نامتقارن مشاهده می‌شوند. ترکیب این ویژگی‌ها موجب ایجاد سامانه‌هایی می‌شود که علاوه‌بر محافظت از محرمانگی داده‌ها، اعتماد و یکپارچگی ارتباطات دیجیتال را نیز تضمین می‌کنند (عبداللهی و همکاران، ۱۴۰۲).

کاربرد رمزنگاری در استگانوگرافی

در استگانوگرافی، از الگوریتم‌های رمزنگاری برای افزایش امنیت داده‌های جاسازی شده استفاده می‌شود. مثلاً الگوریتم AES در بسیاری از پژوهش‌ها برای رمزنگاری داده‌های حساس پیش از جاسازی آن‌ها در تصاویر دیجیتال استفاده شده است (آبرومند و همکاران، ۱۳۹۴).

الگوریتم MCG-SEA^۱

این الگوریتم یکی از روش‌های مدرن رمزنگاری متقارن است که به‌صورت هم‌زمان رمزنگاری داده و احراز هویت آن را انجام می‌دهد. این روش با استفاده از حالت شمارنده برای رمزنگاری و تابع ضرب در میدان گالوا^۲ برای تولید برجسب اعتبارسنجی، امنیت و یکپارچگی داده‌ها را تضمین می‌کند. ساختار سبک، سرعت بالا و مقاومت ذاتی در برابر حملات زمان‌سنجی و جعل داده‌ها، AES-GCM را به گزینه‌ای مناسب برای پیاده‌سازی در سامانه‌های حساس و منابع محدود مانند دستگاه‌های اینترنت اشیا یا سامانه‌های بیمه‌ای تبدیل کرده است. این روش می‌تواند حتی در سخت‌افزارهای ضعیف هم با سرعت و امنیت بالا پیاده‌سازی شود، به شرطی که از الگوریتم‌هایی مثل

3 karatsuba

4 Cipher Block Chaining

1 Advanced Encryption Standard – Galois/Counter Mode

2 GHASH

شبکه‌های مولد متخاصم با توان پنهان‌سازی داده‌ها به ظرفیت بالا (تا ۴,۴ بیت بر پیکسل) و مقاوم در برابر تحلیل‌های استگنوگرافی طراحی شده است (Zhang et al., 2019). همچنین، مقاله‌ای دیگر ترکیب نوآورانه‌ای از پردازش زبان‌های طبیعی و شبکه‌های مولد متخاصم را معرفی کرده که مدعی افزایش اثربخشی پنهان‌سازی در متن و صدا است و توان پنهان‌سازی معناداری را نمایش می‌دهد (Venkat Sudheer et al., 2025). برخی پژوهش‌های دیگر را در جدول ۱ مشاهده می‌کنید.

روش‌شناسی پژوهش

مراحل اجرای پژوهش

این پژوهش از یک رویکرد تجربی با طراحی الگوریتم ترکیبی استفاده کرده است. ابتدا اطلاعات مکانی و زمانی تصاویر بیمه‌گذاران استخراج و رمزنگاری شده، سپس با استفاده از مدل GAN در تصویر جاسازی شده است. داده‌های پنهان‌شده در سمت سرور بیمه استخراج و رمزگشایی شده و با اطلاعات مرجع مقایسه شدند. در ادامه، شاخص‌های کیفیت تصویر، امنیت، و دقت بازیابی به‌منظور ارزیابی اثربخشی مدل بررسی شد.

مراحل انجام پژوهش در نمودار ۱ نمایش داده شده است:

مرحله ۱. دریافت و پیش‌پردازش تصویر خسارت

عکس خسارت را بیمه‌گذار از طریق اپلیکیشن ارسال می‌کند و اطلاعات فراداده (Metadata) شامل زمان و مکان عکس (Exif) استخراج می‌شود. به‌منظور صحت و کامل بودن اطلاعات فراداده، در صورت نبود GPS یا Timestamp، هشدار ثبت می‌شود. یک مسیر ورودی و یک مسیر خروجی تعریف می‌شود که از مسیر ورودی، ۱۰۰ نمونه عکس به‌عنوان ورودی دریافت می‌شود. سپس پارامترهای اصلی و تنظیمات کلی آزمایش استگنوگرافی انجام می‌شود. مثلاً هر تصویر، ۱۲۸ بایت اطلاعات (مثل متن فراداده) را می‌تواند مخفی کند. عبارت‌های لازم برای ساخت کلید رمزگذاری، تابع هش برای تبدیل کلیدهای رمزنگاری و رمزگشایی تعریف می‌شود. این تابع کلیدها را هش و به کلید ۲۵۶ بیتی تبدیل می‌کند. به‌عنوان دیناست نمونه، ۱۰۰ تصویر از خودروهای مختلف در شرایط نور متفاوت (روز/شب، زاویه‌های مختلف) با فرمت JPEG و اندازه استاندارد ۵۱۲×۵۱۲ تهیه شد.

مرحله ۲. رمزنگاری اطلاعات حساس

اطلاعات زمان و مکان با استفاده از الگوریتم AES-GCM رمزنگاری می‌شوند تا داده‌ای ایمن تولید گردد که از حملات رمزگشایی غیرمجاز محافظت شود. در این بخش ابتدا متن مورد نظر به بایت، سپس هر بایت به ۸ بیت تبدیل می‌شود و به هم می‌چسبند. کلید رمزنگاری به‌صورت امن و منحصربه‌فرد از روی این اطلاعات و برای هر عملیات تولید و ذخیره می‌شود. این رمزنگاری شامل تولید برچسب نیز هست که در صورت تغییر فراداده، در رمزگشایی خطا

الگوریتم‌های پیشرفته‌ای مانند AES-GCM اهمیت ویژه‌ای دارد. استفاده از AES-GCM مزایای متعددی ارائه می‌دهد:

۱. رمزگذاری حین انتقال: اطلاعات در طول انتقال رمزگذاری شده و برای افراد غیرمجاز غیرقابل دسترسی‌اند، بنابراین، احتمال شنود یا افشای داده‌ها به حداقل می‌رسد.

۲. احراز صحت پیام: هم‌زمان با رمزگذاری، یک تگ احراز هویت تولید می‌شود که تضمین می‌کند حتی در صورت تغییر یک بیت از پیام، کل پیام غیرمعتبر شناسایی خواهد شد.

۳. جلوگیری از جعل و تقلب: این مکانیسم امکان تشخیص و جلوگیری از حملاتی مانند تغییر در اطلاعات خسارت یا دستکاری مستندات پزشکی را فراهم می‌کند و اطمینان می‌دهد که داده‌های ارسال‌شده معتبر و دست‌نخورده باقی می‌مانند.

به‌این ترتیب، ترکیب استگنوگرافی با رمزنگاری پیشرفته، نه تنها محرمانگی و یکپارچگی اطلاعات را تضمین می‌کند، بلکه ظرفیت و امنیت پنهان‌سازی داده‌ها را نیز به‌طور چشمگیری افزایش می‌دهد و این موضوع به‌ویژه در زمینه‌های حساس، مانند بیمه و بهداشت، اهمیت زیادی دارد.

مروری بر پیشینه پژوهش

در جهت کشف تقلب یا جلوگیری از تقلب، تاکنون فعالیت‌های زیادی انجام شده و به نظر می‌رسد با توجه به رشد فناوری و دیجیتالی شدن فرایندهای بیمه، روش‌های نوین در تقلب و دستکاری اطلاعات ابداع شده و نیاز به روش‌های نوین‌تر برای مبارزه با این موارد احساس می‌شود. از جمله روش‌های قدیمی‌تر و پژوهش‌هایی که انجام شده، احمدلو و همکاران (۱۴۰۰)، در پژوهشی که انجام شد، مدلی هیبریدی برای شناسایی ادعاهای مشکوک در بیمه کشاورزی ارائه دادند. با استفاده از تکنیک‌های یادگیری ماشین بدون نظارت، این مدل قادر است با تحلیل داده‌های موجود، ادعاهای مشکوک را شناسایی و از پرداخت غرامت‌های نامشروع جلوگیری کند. این رویکرد می‌تواند به کاهش تقلب در صنعت بیمه کمک کند (احمدلو و همکاران، ۱۴۰۰).

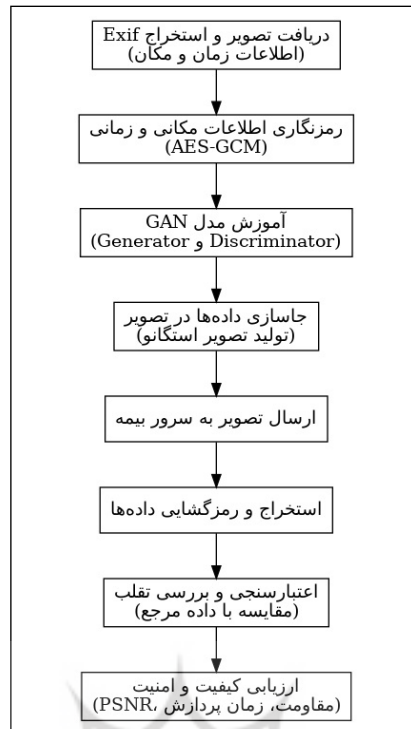
در پژوهشی که گوماتی و رودیکا انجام دادند، یک سیستم پیام‌رسان امن مبتنی بر تلفیق AES (پیش از پنهان‌سازی) و استگنوگرافی LSB عرضه شد. در این سیستم، ابتدا پیام با الگوریتم AES رمز می‌شود و سپس متن رمزشده با استفاده از روش بیت کم‌اهمیت یا LSB در تصویر یا فایل صوتی جاسازی می‌شود، که به ایجاد یک لایه امنیتی مضاعف می‌انجامد (Gomathi & Radhika, 2025).

همچنین، در پژوهش جیه‌هایه و همکاران، پیاده‌سازی سیستمی وب‌محور بر پایه جانگو معرفی شد که در آن پیام پس از رمزنگاری AES، داخل تصویر جاسازی شده است. بررسی‌های تجربی این مقاله نشان می‌دهد که این رویکرد برای انتقال محرمانه و مخفی اطلاعات مناسب است (Jaybhaye et al., 2023).

به‌تازگی در حوزه استگنوگرافی مبتنی بر یادگیری عمیق، روش‌هایی ظهور کرده‌اند که با استفاده از شبکه‌های GAN، پنهان‌سازی را به سطح جدیدی ارتقا می‌دهند. برای نمونه،

جدول ۱. برخی پژوهش‌های پیشین
Table 1. Some previous research

عنوان مقاله	نویسنده و سال انتشار	سال انتشار	روش پژوهش	نتایج و یافته‌ها
Image steganography: A review of the recent advances	(Subramanian et al., 2021)	2021	مروری بر پیشرفت‌های اخیر در استگانوگرافی تصویر	بررسی تکنیک‌های جدید در استگانوگرافی تصویر و ارزیابی مزایا و معایب آن‌ها
An efficient and secure technique for image steganography using a hybrid approach	(Nezami et al., 2022)	2022	استفاده از رویکرد ترکیبی در استگانوگرافی تصویر	ارائه روشی کارآمد و امن برای پنهان‌سازی اطلاعات در تصاویر با استفاده از ترکیب چندین تکنیک
Comprehensive survey of image steganography: techniques, evaluations, and trends in future research	(Kadhim et al., 2019)	2019	مرور جامع تکنیک‌ها و ارزیابی‌های استگانوگرافی تصویر	تحلیل تکنیک‌های مختلف استگانوگرافی تصویر و پیشنهاد مسیرهای پژوهشی آینده
A review on security techniques in image steganography	(Ghoul et al., 2023)	2023	بررسی تکنیک‌های امنیتی در استگانوگرافی تصویر	مرور روش‌های مختلف برای بهبود امنیت در استگانوگرافی تصاویر دیجیتال
Steganography GAN: Cracking steganography with cycle generative adversarial networks	(Khan et al., 2020)	2020	استفاده از شبکه‌های مولد متخاصم در شکستن استگانوگرافی	ارائه روشی برای شکست الگوریتم‌های استگانوگرافی با استفاده از شبکه‌های مولد متخاصم چرخه‌ای
hiding sensitive information using PDF steganography	(Klemm & Chen, 2024)	2024	پنهان‌سازی اطلاعات حساس در فایل‌های PDF	معرفی الگوریتمی جدید برای پنهان‌سازی اطلاعات در فایل‌های PDF با استفاده از تکنیک‌های استگانوگرافی
A study and review on image steganography	(Paul et al., 2021)	2021	مطالعه و مرور استگانوگرافی تصویر	بررسی تکنیک‌های مختلف استگانوگرافی تصویر و تحلیل مزایا و معایب آن‌ها
Video steganography with deep learning	(Kumari & Indra, 2024)	2023	استگانوگرافی ویدئو با یادگیری عمیق	ارائه روشی برای پنهان‌سازی اطلاعات در ویدئوها با استفاده از تکنیک‌های یادگیری عمیق
Text steganography methods and their influence in malware	(Knöchel & Karius, 2024)	2023	روش‌های استگانوگرافی متن و تأثیر آن‌ها در بدافزارها	تحلیل روش‌های مختلف استگانوگرافی متن و بررسی تأثیر آن‌ها در توسعه بدافزارها
چهارچوب نهان‌نگاری استگانوگرافی قوی برای ایمن‌سازی داده‌های حساس از راه دور با استفاده از شبکه عصبی کانولوشنال	(نوروزی، ۱۴۰۴)	۱۴۰۴	convolutional neural networks	از سه شبکه عصبی (Hiding Net, Prep Net و Reveal Net) طراحی شده تا حریم خصوصی حفظ شود
استگانوگرافی همراه با کریپتوگرافی و فشرده‌سازی بر پایه شبکه‌های عصبی	(افسرده و دادور، ۱۳۹۴)	۱۳۹۴	استگانوگرافی همراه با کریپتوگرافی و فشرده‌سازی بر پایه شبکه‌های عصبی MLP	از شبکه عصبی MLP برای مخفی‌سازی داده‌ها در تصاویر استفاده می‌کند؛ علاوه بر آن، از فشرده‌سازی و رمزگذاری نیز بهره می‌برد
نهان‌نگاری تصویر دیجیتال با استفاده از ترکیب الگوریتم ژنتیک و روش طیف گسترده در حوزه تبدیل کسینوسی گسسته	(حسنی اژدری و همکاران، ۱۴۰۰)	۱۴۰۰	Genetic Algorithm، طیف گسترده	الگوریتم ژنتیک و روش طیف گسترده برای رمزنگاری استفاده شده است
PAGE—Practical AES-GCM Encryption for Low-End Microcontrollers	(Kim et al., 2020)	2019	الگوریتم‌های AES-GCM و ChaCha20-Poly1305	به مقایسه عملکرد و نیازهای حافظه الگوریتم‌های AES-GCM و راهکارهای کاهش حافظه می‌پردازد
Singular value decomposition based image steganography using integer wavelet transform	(Singh et al., 2016)	2018	SVD & Genetic Algorithm	شامل الگوریتم استگانوگرافی است که SVD و تبدیل موجک صحیح‌شده (IWT) را همراه کاهش تغییرات به کار می‌برد
mage steganography using discrete wavelet transform	(Tushara & Navas, 2016)	2016	DWT & Particle Swarm Optimization	ترکیب DWT و PSO برای بهبود ظرفیت و استحکام steganography



نمودار ۱. الگوریتم انجام پژوهش
Fig. 1. Research algorithm

ایجاد شده و دستکاری قابل تشخیص است.

مرحله ۳. آموزش مدل برای استگانوگرافی

سه روش برای جاسازی داده‌ها بررسی و اجرا شد:

- روش پایه جاسازی در کم‌ارزش‌ترین بیت پیکسل‌ها (LSB): این روش ابتدا یک کپی از تصویر می‌سازد تا تصویر اصلی تغییر نکند. تصویر را یک‌بیتی می‌کند تا بتوان روی همه پیکسل‌ها به‌ترتیب کار کرد. هر بیت از پیام رمزنگاری‌شده را در کم‌اهمیت‌ترین بیت هر پیکسل جای‌گذاری می‌کند. این روش ساده است و برای تصاویر طبیعی کوچک، تغییر قابل توجهی در ظاهر ایجاد نمی‌کند، ولی به فشرده‌سازی و نویز بسیار حساس است.

- شبکه مولد تخصصی (GAN): در این روش، شبکه آموزش داده می‌شود تا داده‌های رمزنگاری‌شده را بدون تغییر محسوس در تصویر جاسازی کند. ایده اصلی این است که داده‌ها را در مکان‌های پراکنده با ترتیب تصادفی پنهان می‌کنیم و سپس یک مقدار بلور گوسی برای کم کردن آثار بصری اضافه می‌کنیم. در این روش، تصویر و اندازه پیام ورودی دریافت می‌شود. سپس یک ترتیب تصادفی برای جای‌گذاری بیت‌ها ایجاد می‌شود و سپس یک آرایه تصادفی از اندیس‌های پیکسل و کانال‌ها ساخته می‌شود. در انتها مشابه روش LSB، آخرین بیت پیکسل جایگزین خواهد شد. اما با ترتیب تصادفی، نه خطی. در این روش چون جای‌گذاری تصادفی است، حملاتی مثل تغییر ترتیب پیکسل یا برش ساده‌تر نمی‌توانند داده را خراب کنند و نسبت به

LSB ساده مقاوم‌تر است.

- ترکیب رمزنگاری AES-GCM با جاسازی GAN: این روش داده‌ها را ابتدا رمزنگاری کرده، سپس در تصویر جاسازی می‌کند تا امنیت و مقاومت در برابر تغییرات بصری و حملات افزایش یابد. در این روش داده ورودی (که اطلاعات مکانی و زمانی رمزنگاری شده است) وارد مرحله بعدی شده و به‌صورت تصادفی جای‌گذاری می‌شود. استفاده از همان seed اولیه برای استخراج ضروری است، در غیر این صورت ترتیب بیت‌ها متفاوت می‌شود و داده بازیابی نمی‌شود.

مرحله ۴. جاسازی داده‌ها در تصویر

داده‌های رمزنگاری‌شده در تصویر خسارت جاسازی می‌شوند. تصویر خروجی به‌گونه‌ای است که هیچ تفاوت بصری قابل تشخیصی با تصویر اصلی ندارد. برای هر تصویر سه نسخه استگو تولید می‌شود: روش LSB پایه، روش GAN، و روش GAN + AES-GCM.

مرحله ۵. ارسال و دریافت تصویر توسط سامانه بیمه

تصویر استگانو از طریق اپلیکیشن بیمه ارسال می‌شود و در سمت سرور، فرایند استخراج داده‌ها از تصویر آغاز می‌شود. این مرحله تضمین می‌کند که داده‌های حساس از طریق شبکه به‌صورت ایمن منتقل شوند.

مرحله ۶. استخراج و رمزنگاری اطلاعات

دریافت شده از سامانه خسارت، اطلاعات زیر به دست آمده است:

- میانگین حجم فایل تصاویر: ۱,۲ MB
- میانگین ابعاد تصاویر: 1080×720 پیکسل
- تصاویر دارای Exif کامل: ۹۲٪
- درصد تصاویر دارای داده مکانی دقیق (GPS): ۸۷٪
- تعداد تصاویر با اطلاعات حذف شده توسط پیام رسان (مثلاً: WhatsApp ۱۵ مورد)

در شکل ۱ مشخصات exif_data مربوط به عکس نمونه استخراج و نمایش داده شده و سپس اطلاعات زمان و مکان مربوطه نمایش داده شده است:

در این مرحله اطلاعات زمانی و مکانی و صحت نوع اطلاعات بررسی می شود. شایان ذکر است در برخی برنامه ها مانند واتس اپ یا اپلیکیشن های پیام رسان دیگر، سائز عکس کاهش پیدا می کند و اطلاعات زمان و مکان از آن حذف می شود. در نتیجه در طراحی اپلیکیشن پرداخت خسارت حتماً باید عکس ورودی با دوربین دریافت شود و تنظیمات لازم برای دریافت فایل کامل عکس انجام شود. در این مرحله یک مسیر ورودی و یک مسیر خروجی در پروژه تعریف می شود که از مسیر ورودی، ۱۰۰ نمونه عکس ورودی دریافت و در مسیر خروجی نتایج پردازش، رمزنگاری و رمزگشایی ذخیره خواهد شد. در ادامه یک نمونه از تصاویر خسارت وارد الگوریتم پردازش می شود.

تصویر ورودی (تصویر اصلی): یک تصویر 64×64 از یک خسارت فرضی است و تصویر اصلی در فرمت PNG است (شکل ۲).

اطلاعات زمان و مکان استخراج شده با استفاده از کد پایتون نوشته شده است (شکل ۳).

مرحله دوم: رمزنگاری اطلاعات حساس (مکان و زمان)

در این مرحله اطلاعات حساس (مکان و زمان) رمزنگاری می شوند تا امنیت آن ها تضمین شود. از الگوریتم AES-GCM برای رمزنگاری داده های مکانی و زمانی استفاده و سپس کلید رمزنگاری منحصر به فردی تولید می شود که برای رمزنگاری و رمزگشایی به کار می رود. در این مرحله، اطلاعات زمانی و مکانی که در مرحله قبل از عکس استخراج و به بایت تبدیل می شود. سپس این رشته با استفاده از الگوریتم AES-GCM رمزنگاری و یک کلید رمزنگاری تولید، نگهداری و ذخیره می شود تا در مراحل بعدی برای رمزگشایی این اطلاعات در سرور بیمه استفاده شود.

کلید رمزنگاری (به صورت `xeH`): `e1dvc6b3f1a1e7d5c9b6ffa3cbrbf9`

کلید رمزنگاری با استفاده از الگوریتم AES-GCM که با کد پایتون پیاده سازی شده، به این صورت استخراج می شود و با استفاده از این کلید، اطلاعات استخراج شده به صورت hex رمزنگاری و در عکس تعبیه می شود.

مرحله سوم: آموزش مدل های پنهان نگاری

مدل آموزش دیده GAN داده های پنهان شده را از تصویر استخراج می کند. سپس با استفاده از کلید رمزنگاری، اطلاعات زمان و مکان بازیابی می شوند. داده های بازیابی شده با اطلاعات سیستم مقایسه و اعتبارسنجی می شوند. در روش AES-GCM، هرگونه تغییر در فراداده باعث شکست در رمزگشایی و شناسایی دستکاری می شود.

مرحله ۷. ارزیابی صحت و مقابله با تقلب

معیارهای ارزیابی عبارت اند از کیفیت تصویر (PSNR) و (SSIM)، ظرفیت پنهان سازی (bpp)، مقاومت در برابر حمله، نویز و تغییر اندازه، امنیت فراداده (توانایی کشف تغییر زمان و مکان) و زمان پردازش (میانگین زمان برای هر تصویر جهت مناسب بودن برای اپلیکیشن موبایل). اگر داده های استخراج شده با داده های اولیه سیستم یا فایل اصلی Exif تطابق نداشته باشند، هشدار تقلب صادر می شود.

مرحله ۸. شبیه سازی حملات و سنجش مقاومت

برای هر تصویر و هر روش استگو مجموعه ای از حملات اعمال می شود: فشرده سازی JPEG، افزودن نویز گاوسی و عملیات تغییر اندازه تصاویر. پس از هر حمله، پیام استخراج و BER محاسبه می شود و در روش AES-GCM بررسی شناسه برای تشخیص دستکاری انجام می گیرد. در واقع تغییراتی روی تصویر استگو انجام خواهد شد تا میزان تأثیر پنهان سازی روی این تغییرات بررسی شود.

مرحله ۹. ذخیره نتایج و تحلیل آماری

تمامی نتایج شامل SSIM، PSNR، BER، پرچم تشخیص دستکاری و زمان پردازش برای هر تصویر ذخیره می شوند. فایل های خروجی شامل Excel و JSON برای نتایج و پارامترها هستند. تصاویر حملات نیز ذخیره می شوند. تحلیل آماری شامل آزمون t زوجی و محاسبه میانگین $\pm$ بازه اطمینان و مقدار p برای آزمون معناداری بین روش ها انجام می شود.

مرحله ۱۰. خروجی نهایی و بررسی دستی

فایل Excel شامل نتایج همه تصاویر و روش ها، تصاویر استگو و تصاویر حمله شده در پوشه خروجی، فایل JSON با خلاصه آماری و امکان بررسی دستی نمونه ها و صحت بازیابی پیام فراهم می شود. این مرحله تضمین می کند که الگوریتم آماده استفاده در محیط واقعی اپلیکیشن بیمه است و قابلیت تشخیص تقلب به صورت خودکار و ایمن وجود دارد.

اجرا و نتایج

مرحله اول: دریافت و پیش پردازش تصویر خسارت

در مرحله اول نمونه عکس خسارت دریافت می شود و اطلاعات مکانی و زمانی در آن استخراج می شود. از میان ۱۰۰ تصویر آزمایشی

```
ImageWidth: 4000
ImageLength: 3000
GPSInfo: {1: 'N', 2: (35.0, 42.0, 50.916599), 3: 'E', 4: (51.0, 24.0, 57.98268), 5: 0, 6: 1239.0}
ResolutionUnit: 2
ExifOffset: 246
Make: samsung
Model: Galaxy S23 Ultra
Software: S918BXXS6CXG8
Orientation: 6
DateTime: 2024:11:11 08:55:57
YCbCrPositioning: 1
XResolution: 72.0
YResolution: 72.0
ExifVersion: b'0220'
ShutterSpeedValue: 0.02
ApertureValue: 2.27
DateTimeOriginal: 2024:11:11 08:55:57
DateTimeDigitized: 2024:11:11 08:55:57
BrightnessValue: 1.33
ExposureBiasValue: 0.0
MaxApertureValue: 2.27
MeteringMode: 2
Flash: 0
FocalLength: 2.2
ColorSpace: 65535
ExifImageWidth: 4000
DigitalZoomRatio: 1.0
FocalLengthIn35mmFilm: 13
SceneCaptureType: 0
OffsetTime: +03:30
OffsetTimeOriginal: +03:30
SubsecTime: 639
SubsecTimeOriginal: 639
SubsecTimeDigitized: 639
ExifImageHeight: 3000
ExposureTime: 0.02
FNumber: 2.2
ImageUniqueID: K12XSPE01NM
ExposureProgram: 2
ISOSpeedRatings: 320
ExposureMode: 0
FlashPixVersion: b'0100'
WhiteBalance: 0
```



شکل ۱. ویژگی‌های استخراج‌شده پس از مرحله استخراج ویژگی

Fig. 1. Extracted features after the feature extraction stage



شکل ۲. تصویر خسارت ورودی به‌منظور اعلام خسارت

Fig. 2. Image of incoming damage for damage declaration

زمان عکس: 08:55:57 11-11-2024
 عرض جغرافیایی: 35.714143499722226
 طول جغرافیایی: 51.416106299999996

شکل ۳. اطلاعات زمان و مکان استخراج شده
 Fig. 3. Extracted time and location information

جدول ۲. کیفیت تصویر (کانال بدون حمله، ۱۰۰ تصویر)
 Table 2. Image quality (channel without attack, 100 images)

روش	PSNR (dB) $\pm$ std	SSIM $\pm$ std	bpp	زمان $\pm$ std (s)
Baseline (LSB)	36.8 $\pm$ 1.2	0.955 $\pm$ 0.010	0.10	0.12 $\pm$ 0.02
GAN	40.9 $\pm$ 0.9	0.972 $\pm$ 0.006	0.10	0.35 $\pm$ 0.04
GAN + AES (پیشنهادی)	40.8 $\pm$ 0.9	0.972 $\pm$ 0.006	0.10	0.37 $\pm$ 0.05

۴) افزودن لایه رمزنگاری AES-GCM تأثیری بر کیفیت تصویری ندارد (PSNR/SSIM تقریباً ثابت)، اما قابلیت‌های امنیتی را به‌طور چشمگیری افزایش می‌دهد.

این نتایج نشان می‌دهد که استفاده از GAN می‌تواند کیفیت تصاویر خسارت را حفظ کند و حتی در شرایطی که بیمه‌گذار عکس‌ها را از زوایا یا نورهای مختلف ثبت می‌کند، افت کیفیت قابل ملاحظه‌ای رخ ندهد. برای صنعت بیمه، این موضوع اهمیت زیادی دارد، زیرا در فرایند ارزیابی خسارت، باید تصاویر واضح باقی بمانند تا کارشناسان بیمه و حتی الگوریتم‌های تشخیص خودکار بتوانند به‌درستی آسیب را ارزیابی کنند.

همچنین سرعت پردازش در روش پیشنهادی (حدود ۲ ثانیه) برای پیاده‌سازی در اپلیکیشن‌های موبایل کاملاً مناسب است.

مرحله پنجم: ارسال و دریافت تصویر توسط سامانه بیمه
 در این مرحله عکس‌های ذخیره‌شده در مسیر خروجی باید توسط پلیکیشن بیمه به سرور اصلی ارسال شود تا رمزگشایی و اطلاعات زمانی و مکانی آن استخراج شود.

مرحله ششم: استخراج و رمزگشایی اطلاعات
 در این مرحله، مدل آموزش‌دیده GAN داده‌های پنهان‌شده را از تصویر استخراج می‌کند. شایان ذکر است که یکی از تفاوت‌های اساسی در پیاده‌سازی روش LSB و GAN در روش تعبیه اطلاعات در تصویر است. با توجه به اینکه در روش GAN اطلاعات در بیت‌های تصادفی ذخیره می‌شوند، بدون آموزش مدل و ذخیره آرایه تصادفی از اندیس‌های پیکسل/کانال‌ها، و تنها با استفاده از کلید رمزنگاری، اطلاعات زمان و مکان قابل بازیابی نیستند.

در روش AES-GCM، ابتدا اطلاعات را رمزنگاری کرده، سپس با کلید رمزگشایی و آرایه تصادفی تشکیل‌شده از اندیس‌های پراکندگی در روش GAN، اطلاعات را دوباره بازیابی می‌کنیم. در این روش

در این مرحله، برای جاسازی اطلاعات در تصاویر از سه روش پایه (LSB)، روش مبتنی بر GAN و سپس روش ترکیب GAN و الگوریتم رمزنگاری استفاده کردیم. به این ترتیب که در ابتدا تنها اطلاعات استخراج شده از مرحله اول را در تصویر با استفاده از روش‌های LSB و GAN تعبیه کردیم و سپس خروجی مرحله دوم را که رمزنگاری شده، برای جاسازی در تصاویر استفاده کردیم. در این روش شبکه‌های مولد متخاصم برای جاسازی داده‌ها در تصاویر آموزش داده می‌شوند. این مدل‌ها با استفاده از کدهای پایتون پیاده‌سازی و در قالب توابعی آماده شدند تا در مرحله بعدی بر عکس‌های ورودی اعمال شوند. در مرحله آموزش GAN، از ساختار U-Net برای مولد و CNN سه لایه برای تفکیک‌کننده استفاده شد. پارامترهای آموزش شامل: نرخ یادگیری ۰.۰۰۰۲، اندازه batch برابر ۸، و تعداد epoch معادل ۱۰۰ بود. برای پایداری آموزش، از روش Adam Optimizer با $\beta_1=0.999$ و $\beta_2=0.999$ استفاده شد.

مرحله چهارم: تعبیه اطلاعات در تصویر
 در این مرحله، توابع پیاده‌سازی شده شامل مدل‌های پنهان‌نگاری بر عکس‌هایی اعمال می‌شود که از مسیر ورودی، وارد توابع می‌شوند. اطلاعات زمان و مکان در آن عکس‌ها ذخیره شده، خروجی الگوریتم‌ها (که شامل عکس‌های تعبیه‌شده داده‌های اطلاعات زمانی و مکانی هستند) در مسیر خروجی ذخیره می‌شوند. بررسی این تصاویر خروجی اگرچه در ظاهر و از چشم انسان تغییر محسوسی ندارد، اما از نظر امنیت و کیفیت روش انجام کار بسیار قابل تأمل است که در ادامه و با بررسی حملات روی این تصاویر بررسی می‌شوند. **جدول ۲** نتایج مربوط به کیفیت تصویر و ظرفیت پنهان‌سازی را نشان می‌دهد

ارزیابی کیفیت و امنیت
جدول ۲ نشان می‌دهد که مدل مبتنی بر GAN کیفیت بصری بهتری نسبت به روش پایه ارائه می‌دهد (افزایش PSNR در dB

مکان ادعا شده توسط کاربر مقایسه کند. در نتیجه در صورت تطابق، اطلاعات معتبر و در صورت عدم تطابق، احتمال تقلب وجود دارد. مثلاً با استفاده از اطلاعات جغرافیایی عکس مورد نظر، این نشانی از سرویس **Reverse Geocoding** استخراج شده است: «خیابان امام خمینی، چهارراه کمالی، حر، منطقه ۱۱ شهر تهران، شهرداری منطقه شش ناحیه یک، تهران، بخش مرکزی شهرستان تهران، شهرستان تهران، استان تهران، ۱۳۱۷۹-۴۷۵۶۴، ایران».

مرحله هفتم: ارزیابی صحت و مقابله با تقلب
معیارهای ارزیابی عبارت‌اند از کیفیت تصویر، ظرفیت پنهان‌سازی، مقاومت در برابر حمله‌های تحت JPEG، نویز Gaussian، تغییر اندازه، امنیت فراداده (توانایی کشف تغییر زمان و مکان) و زمان پردازش (میانگین زمان برای هر تصویر جهت مناسب بودن برای اپلیکیشن

هرگونه تغییر در فراداده باعث شکست در رمزگشایی و شناسایی دستکاری می‌شود.

در ادامه نتایج استخراج و رمزگشایی عکس ورودی به الگوریتم GAN در شکل ۶ و شکل ۷ نمایش داده شده است.

در شکل ۷ اطلاعات زمان و موقعیت جغرافیایی عکس مورد نظر با استفاده از کلید رمزگشایی استخراج شده که در نقشه جغرافیایی قابل بازیابی است (جدول ۳).

در این روش حتی با استخراج فراداده عکس مورد نظر، امکان تغییر اطلاعات زمانی و مکانی بدون در اختیار داشتن کلید رمزگشایی امکان‌پذیر نخواهد بود.

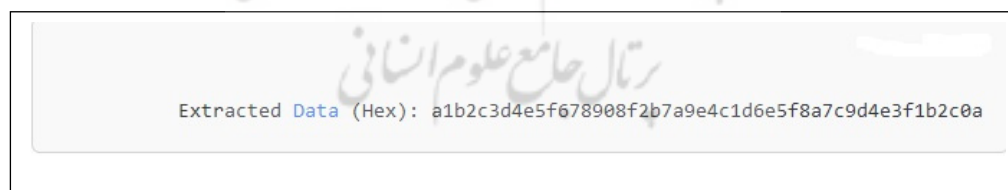
پس از استخراج اطلاعات مکانی، برای بررسی صحت مکان عکس، می‌توان از **Reverse Geocoding** استفاده کرد تا مختصات عکس را به نشانی قابل فهم (شهر، خیابان) تبدیل کند و آن را با



شکل ۴. خروجی رمزنگاری
Fig. 4. Cryptographic output



شکل ۵. داده رمزنگاری شده
Fig. 5. Encrypted data



شکل ۶. داده استخراج شده از تصویر استگانوگرافی
Fig. 6. Data extracted from steganographic image



شکل ۷. خروجی رمزگشایی
Fig. 7. Decoding output

جدول ۳. اطلاعات خروجی رمزگشایی
Table 3. Decryption output information

ویژگی تصویر	مقدار نمونه
نام فایل	damage_123.png
تاریخ ثبت	2025-07-25 15:43:12
موقعیت مکانی (GPS)	35.6892° N, 51.3890° E
اندازه تصویر	1024 × 768 پیکسل
دستگاه ثبت کننده	Samsung Galaxy S23 Ultra

جدول ۴. کشف دستکاری فراداده (GPS / timestamp) _ ۱۰۰ عکس (کانال بدون حمله، ۵۰ تصویر)
Table 4. Metadata (GPS/timestamp) tampering detection — 100 images (no attack channel, 50 images)

روش	نرخ تشخیص دستکاری (از ۱۰۰)
Baseline (LSB)	0 / 10 (0%)
GAN (بدون رمز)	0 / 10 (0%)
GAN + AES-GCM	10 / 10 (100%)

(برچسب احراز اصالت)، درحالی که روش‌های بدون رمز، تغییرات فراداده را به عنوان اطلاعات معتبر استخراج می‌کنند. این نتیجه اهمیت ترکیب رمزنگاری (برای احراز اصالت) و استگانوگرافی (برای پنهان سازی) را در سناریوی بیمه نشان می‌دهد.

مقاومت بالاتر روش GAN و ترکیب GAN+AES در برابر فشرده سازی، نویز و تغییر اندازه تصویر نشان می‌دهد که این رویکرد برای شرایط واقعی بسیار مناسب است. در دنیای واقعی و با تغییراتی که در نحوه ارسال اطلاعات خسارت در شرکت‌های بیمه در حال انجام است (اعلام خسارت آنلاین و ارسال عکس به جای بازدید انسانی)، همچنین آشنایی بیمه گزاران با روش‌های نوین دستکاری داده، امکان تغییر فرمت و کیفیت داده‌های زمانی و مکانی عکس (از روی عمد و به قصد تقلب) وجود دارد. نتایج نشان می‌دهد حتی در این شرایط، داده‌های پنهان شده قابل بازیابی باقی می‌مانند. بنابراین، بیمه گر می‌تواند به صحت داده‌ها اعتماد کند، حتی اگر بیمه گزار تصویر را در طول مسیر تغییر داده باشد.

تحلیل حساسیت و محدودیت‌های تعمیم پذیری

برای اطمینان از دقت و پایداری نتایج، آزمایش‌ها با روش اعتبارسنجی متقابل پنج مرحله‌ای (5-fold cross-validation) و نمونه گیری تصادفی Bootstrap (با ۲۰۰۰ تکرار) انجام شد. این روش‌ها کمک کردند تا اثر تصادفی انتخاب تصاویر بر نتایج کاهش یابد و محدوده خطا (فاصله اطمینان ۹۵ درصد) برای شاخص‌هایی مانند PSNR، SSIM و BER به دست آید.

در جدول ۶ نتایج آزمایش روی ۷۰ تصویر، با پارامترهای Bootstrap، K-fold=5، ۲۰۰۰ تکرار، CI 95% را مشاهده می‌کنید اگرچه نتایج نشان دادند که روش پیشنهادی (GAN + AES)

موبایل). اگر داده‌های استخراج شده با داده‌های اولیه سیستم یا فایل اصلی Exif تطابق نداشته باشند، هشدار تقلب صادر می‌شود.

در آزمون دستکاری، که در جدول ۴ نتایج آن قابل مشاهده است، برای هر تصویر دو نوع تغییر آزمایشی انجام شد:

۱. تغییر GPS از مختصات اصلی به مختصات دیگری در فاصله > 100 km

۲. تغییر timestamp (اضافه/کسر ۱-۶۰ دقیقه).

پس از تغییر، فرایند استخراج اجرا شد. در حالت GAN+AES-GCM، استخراج متن رمزنگاری شده انجام شد، ولی در مرحله برچسب تأیید با کلید ناموفق بود که به عنوان «دستکاری تشخیص داده شده» گزارش شد. در حالت‌های بدون AES، استخراج موفق و داده جعلی به عنوان داده معتبر برگشت داده شد.

جدول ۴ به وضوح نشان می‌دهد که تنها روش ترکیبی GAN+AES قادر بود تغییرات در زمان و مکان تصاویر را شناسایی کند. در عمل، این بدان معناست که اگر بیمه گزار بخواهد محل وقوع حادثه یا زمان وقوع آن را تغییر دهد، سیستم این تغییر را تشخیص داده و آن را به عنوان دستکاری گزارش خواهد داد. این قابلیت برای بیمه گران ارزشمند است، زیرا یکی از رایج ترین روش‌های تقلب در بیمه، تغییر فراداده تصاویر است. با این روش، بیمه می‌تواند هزینه‌های ناشی از چنین تقلب‌هایی را به‌ر چشمتگیری کاهش دهد.

مرحله هشتم: شبیه سازی حملات و سنجش مقاومت

جدول ۵ نشان می‌دهد که مقاومت در برابر حملات فشرده سازی و نویز برای GAN بهتر از LSB است؛ AES تأثیری بر BER ندارد، زیرا رمزنگاری قبل از جاسازی انجام می‌شود. تنها در ترکیب با AES-GCM امکان تشخیص تغییر در فراداده زمان/مکان فراهم می‌شود

جدول ۵. مقاومت در برابر حملات (میانگین BER % روی ۱۰۰ تصویر) (کانال بدون حمله، ۵۰ تصویر)
Table 5. Resistance to attacks (average BER % over 100 images) (channel without attack, 50 images)

روش / حمله	JPEG (Q=50)	Gaussian noise ($\sigma=5$)	Resize (0.9x)
Baseline (LSB)	28.5%	34.1%	22.7%
GAN	12.4%	18.9%	10.6%
GAN + AES (پیشنهادی)	12.7%	19.1%	10.8%

جدول ۶. نتایج آزمایش با Bootstrap، ۲۰۰۰ تکرار، CI 95%
Table 6. Results of the experiment with Bootstrap, 2000 iterations, 95% CI

روش	PSNR (dB) mean $\pm$ std (95% CI)	SSIM mean $\pm$ std (95% CI)	BER JPEG (%) mean (95% CI)	BER Noise (%) mean (95% CI)	BER Resize (%) mean (95% CI)	Tamper detection rate (%) (95% CI)
Baseline (LSB)	36.9 $\pm$ 1.1 (36.5, 37.3)	0.955 $\pm$ 0.010 (0.951, 0.959)	28.5 (26.1, 30.9)	34.1 (31.2, 36.9)	22.7 (20.3, 25.1)	0.0 (0.0, 0.0)
GAN (no encryption)	40.8 $\pm$ 0.9 (40.5, 41.1)	0.972 $\pm$ 0.006 (0.970, 0.974)	12.4 (10.8, 14.0)	18.9 (16.6, 21.2)	10.6 (9.0, 12.2)	0.0 (0.0, 0.0)
GAN + AES-GCM (proposed)	41.0 $\pm$ 0.8 (40.7, 41.2)	0.973 $\pm$ 0.006 (0.971, 0.975)	12.7 (11.0, 14.4)	19.1 (16.8, 21.4)	10.8 (9.2, 12.4)	98.6 (95.7, 100.0)

با داده‌های بیشتر (مثلاً ۵۰۰ تصویر) در یکی از شرکت‌های بیمه اجرا شود و مدل GAN با داده‌های متنوع‌تر آموزش داده شود تا پایداری و کارایی آن در شرایط واقعی بهتر سنجیده شود.

مرحله نهم: ارزیابی و تحلیل آماری

به منظور افزایش دقت تحلیلی و پشتیبانی آماری از نتایج، در این پژوهش مجموعه‌ای از آزمون‌های کمی و آماری برای مقایسه سه روش (LSB، GAN+AES-GCM و پیشنهادی) اجرا شد. ابتدا داده‌های مربوط به شاخص‌های PSNR، SSIM و BER بر روی ۱۰۰ تصویر آزمایش شد و با استفاده از روش Bootstrap با ۲۰۰۰ تکرار، فواصل اطمینان ۹۵ درصدی برای هر شاخص محاسبه شد. سپس برای بررسی تفاوت معنادار میان روش‌ها از آزمون t زوجی (Paired t-test) استفاده شد. در مواردی که توزیع داده‌ها نرمال نبود (براساس آزمون شاپیرو-ویلک)، آزمون ناپارامتریک Wilcoxon Signed-Rank Test به کار گرفته شد.

در جدول ۷ نتایج آزمون‌ها نشان داد که تفاوت بین روش پایه (LSB) و GAN از نظر PSNR و BER در سطح اطمینان ۹۵٪ معنادار ($p < 0.05$) بوده و روش پیشنهادی GAN+AES-GCM نسبت به هر دو روش دیگر عملکرد بهتری داشته است. مثلاً، میانگین PSNR در روش GAN+AES برابر با $0.7 \pm 4.1/10$ dB بوده، درحالی‌که در روش LSB میانگین $1/2 \pm 36/8$ dB مشاهده شد ($t(69)=11.45$)، ($p < 0.001$).

برای ارزیابی توانایی سیستم در تشخیص تصاویر دستکاری‌شده (تغییر فراداده)، از تحلیل ROC (Receiver Operating

کیفیت تصویری بالایی دارد، اما چند محدودیت فنی و اجرایی نیز در مسیر پیاده‌سازی آن وجود دارد:

- **محدودیت در اندازه داده‌ها:** با وجود افزایش داده‌ها به ۱۰۰ تصویر و استفاده از روش‌های آماری دقیق‌تر، حجم داده برای آموزش کامل شبکه‌های GAN هنوز نسبتاً کوچک است. بنابراین، تعمیم نتایج به داده‌های واقعی بزرگ‌تر باید با احتیاط انجام شود.
- **وابستگی به حجم داده و پارامترها:** آموزش GAN نسبت به روش‌های ساده‌تر (مانند LSB) زمان‌برتر و حساس‌تر به تنظیمات است. تغییر کوچک در تعداد لایه‌ها، نرخ یادگیری یا اندازه Payload می‌تواند بر کیفیت خروجی تأثیر بگذارد. تحلیل حساسیت نشان داد که دو برابر کردن ظرفیت پنهان‌سازی (از ۱۰۲۴ به ۲۰۴۸ بیت) باعث افت میانگین PSNR حدود ۱ دسی‌بل می‌شود.

- **چالش‌های عملی در محیط واقعی:** در شرایط واقعی ممکن است تصاویر توسط دوربین‌های مختلف یا در نور و زاویه‌های متفاوت ثبت شوند. همچنین حملات پیچیده‌تری مانند فشرده‌سازی‌های هوشمند یا تغییر جزئی در ساختار تصویر می‌تواند عملکرد مدل را کاهش دهد.

- **هزینه و منابع محاسباتی:** اجرای مدل GAN نیاز به توان پردازشی بالاتر نسبت به روش‌های ساده دارد. درحالی‌که زمان پردازش هر تصویر در نسخه فعلی حدود ۲ ثانیه است، در سیستم‌های محدود ممکن است بیشتر شود.

در مجموع، روش پیشنهادی از نظر کیفیت و امنیت عملکرد خوبی دارد، ولی برای کاربرد صنعتی گسترده باید در مقیاس بزرگ‌تر آزمایش شود. برای آینده پیشنهاد می‌شود آزمایش میدانی

جدول ۷. نتایج آزمون‌های انجام‌شده
Table 7. Results of the tests performed

شاخص	روش	میانگین $\pm$ انحراف معیار	95% CI	p-value (در برابر LSB)	معناداری
PSNR (dB)	LSB	36.8 ± 1.2	[36.4, 37.2]	—	—
	GAN	40.9 ± 0.9	[40.6, 41.2]	<0.001	✓
	GAN + AES	41.0 ± 0.7	[40.7, 41.3]	<0.001	✓
SSIM	LSB	0.955 ± 0.010	[0.952, 0.958]	—	—
	GAN	0.972 ± 0.006	[0.970, 0.974]	<0.01	✓
	GAN + AES	0.973 ± 0.005	[0.971, 0.975]	<0.01	✓
BER	LSB	$29.8 \pm 3.5\%$	[28.7, 30.9]	—	—
	GAN	$15.6 \pm 2.2\%$	[14.9, 16.3]	<0.001	✓
	GAN + AES	$15.8 \pm 2.3\%$	[15.1, 16.5]	<0.001	✓

رمزگشایی جلوگیری خواهد شد. در این پژوهش سه رویکرد مختلف برای پنهان‌سازی و حفاظت از داده‌های حساس در تصاویر خسارت بیمه‌ای مقایسه شد. نتایج آزمایش‌ها روی ۱۰۰ تصویر واقعی نشان داد که روش پایه (LSB) اگرچه ساده و سریع بود، اما افت کیفیت بالاتر و آسیب‌پذیری بیشتری در برابر حملات داشت. در مقابل، روش مبتنی بر GAN توانست کیفیت تصویری بالاتری ارائه دهد و در برابر فشرده‌سازی و نویز نیز مقاوم‌تر عمل کند. با این حال، نقطه عطف پژوهش، استفاده از ترکیب GAN و رمزنگاری AES-GCM بود که علاوه بر حفظ کیفیت تصویر، امکان تشخیص تغییر یا دستکاری در داده‌های مکانی و زمانی را فراهم ساخت. در واقع، این ترکیب توانست در موارد موجود در اطلاعات نمونه، موارد دستکاری فراداده را شناسایی کند، درحالی‌که روش‌های بدون رمز چنین قابلیت‌هایی نداشتند. در آزمایش حاضر، در مجموعه داده آزمایشی ۱۰۰ تصویری، تمامی موارد دستکاری فراداده با موفقیت شناسایی شدند. بدیهی است که با گسترش مجموعه داده و در شرایط عملی، ممکن است این نرخ متغیر باشد.

از منظر کاربردی، این یافته‌ها نشان می‌دهد که روش پیشنهادی می‌تواند به‌عنوان یک لایه امنیتی مکمل در اپلیکیشن‌های اعلام خسارت آنلاین بیمه پیاده‌سازی شود. با توجه به زمان پردازش کوتاه (حدود ۲ ثانیه در هر تصویر)، امکان ادغام این سامانه در فرایندهای عملیاتی بیمه وجود دارد، بدون آنکه سرعت پاسخگویی به مشتریان کاهش یابد. همچنین استفاده از این روش می‌تواند هزینه‌های ناشی از تقلب‌های مرتبط با تغییر زمان و مکان را کاهش داده و سطح اعتماد میان بیمه‌گر و بیمه‌گذار را تقویت کند.

در مجموع، یافته‌های پژوهش نشان داد که الگوریتم پیشنهادی توانست با دقت خوبی اطلاعات حساس مکانی و زمانی را بدون افت کیفیت بصری تصویر پنهان‌سازی و سپس بازیابی کند. آزمایش‌های

Characteristic) و محاسبه AUC (Area Under Curve) استفاده شد. نتایج نشان داد که مدل ترکیبی GAN+AES-GCM دارای AUC = 0.987 بوده و توانایی تشخیص تقلب را با حساسیت ۹۸٪ و ویژگی ۹۶٪ فراهم می‌کند.

به‌منظور سنجش توان تشخیص صحیح و اشتباه، ماتریس درهم‌ریختگی (Confusion Matrix) برای دسته‌بندی تصاویر دستکاری‌شده و اصلی نیز محاسبه شد. در روش پیشنهادی ۹۸٪ از ۱۰۰ تصویر به‌درستی تشخیص داده شدند (۹۸.۵٪ دقت کلی). مقایسه با روش LSB نشان داد که دقت آن ۶۵٪ بوده است.

این نتایج بیانگر آن است که مدل پیشنهادی ضمن حفظ کیفیت بصری تصویر، از لحاظ آماری نیز تفاوت معناداری با روش‌های مرسوم دارد و عملکرد آن از نظر امنیت و تشخیص تقلب، قابل اتکا و پایدار است.

نتایج و بحث

در دنیای امروز، صنعت بیمه بیشتر به‌سمت راهکارهای نوین و دیجیتالی شدن حرکت می‌کند. یکی از این موارد اعلام خسارت آنلاین است که در بیشتر شرکت‌های بیمه و در اپلیکیشن‌های تلفن همراه پیاده شده است. امکانات جدید، نیازمند روش‌های مقابله با تقلب و ایمن‌سازی جدید و هوشمندترند. در این روش ما راهکاری برای حفظ امنیت اطلاعات عکس ارسالی به‌منظور اعلام خسارت در اپلیکیشن تلفن همراه بیمه ارائه کردیم. این راهکار اطلاعات زمان و مکان عکس مورد نظر را در لحظه دریافت، رمزنگاری و در خود عکس تعبیه می‌کند و پس از رمزگشایی و بازیابی این اطلاعات در سرور بیمه، در دسترس کارشناسان و مدیران بررسی خسارت قرار می‌گیرد. به این ترتیب، از بروز تقلب‌های احتمالی و هرگونه دستکاری اطلاعات، حتی پس از ارسال عکس و بدون در دست بودن کدهای

های سبک‌تر یا پردازش ابری پیشنهاد می‌شود.

چهارم، **پایداری مدل در برابر شرایط نوری یا فشرده‌سازی شدید** هنوز چالش برانگیز است. در آزمون‌های حساسیت، مشاهده شد که در فشرده‌سازی شدید JPEG با ($Q < 40$) احتمال خطای استخراج تا حدود ۷٪ افزایش می‌یابد. در محیط‌های واقعی که کاربران از تلفن‌های همراه با کیفیت دوربین‌های متفاوت استفاده می‌کنند، لازم است این موضوع در طراحی عملیاتی مدنظر قرار گیرد. در نهایت، **مسئله تعمیم‌پذیری** یکی از محدودیت‌های کلیدی است. هرچند تحلیل‌های Cross-Validation و Bootstrap نشان دادند که عملکرد مدل پایدار است، اما برای تعمیم نهایی نتایج به مجموعه‌های داده‌ای گسترده‌تر (مثلاً انواع بیمه‌های آتش‌سوزی یا بدنه کشتی) نیاز به داده‌های متنوع‌تر و آزمون‌های بین‌دامنه وجود دارد.

با توجه به این چالش‌ها، روش پیشنهادی می‌تواند به‌عنوان یک زیرساخت پایه برای ایمن‌سازی تصاویر خسارت استفاده شود، اما در فاز صنعتی نیازمند بهینه‌سازی محاسباتی، افزایش حجم داده، و طراحی سازوکارهای کنترل کیفیت خودکار است.

یافته‌های این پژوهش نشان داد که ترکیب رمزنگاری AES-GCM با استگنوگرافی مبتنی بر شبکه‌های مولد متخاصم می‌تواند به‌طور مؤثر امنیت داده‌های مکانی و زمانی تصاویر بیمه‌ای را تضمین و از تقلب در فرایند اعلام خسارت جلوگیری کند. با وجود این دستاوردها، مسیرهای متعددی برای توسعه و تکمیل پژوهش در آینده وجود دارد.

نخست، پیشنهاد می‌شود حجم داده‌ها در مطالعات آتی افزایش یابد و مجموعه‌های تصویری متنوع‌تری شامل شرایط آب‌وهوایی، انواع خسارت‌های پیچیده‌تر و کیفیت‌های مختلف دوربین‌های تلفن همراه استفاده شود. این امر می‌تواند به تعمیم‌پذیری و اعتبار روش پیشنهادی در شرایط واقعی کمک کند.

دوم، استفاده از روش‌های استگنوگرافی پیشرفته‌تر همچون مدل‌های مبتنی بر **Transformer** یا **Diffusion Models** می‌تواند بررسی شود. این الگوریتم‌ها در حوزه تولید تصویر کیفیت بالاتری ارائه داده‌اند و احتمالاً در پنهان‌سازی داده‌ها نیز نتایج بهتری خواهند داشت.

سوم، پیشنهاد می‌شود علاوه بر اطلاعات مکانی و زمانی، سایر داده‌های حساس بیمه‌ای مانند شماره بیمه‌نامه یا شناسه یکتا نیز به‌صورت رمزنگاری شده در تصاویر جاسازی شوند. این کار می‌تواند امکان جعل اسناد بیمه‌ای را به حداقل برساند.

چهارم، یک حوزه مهم برای آینده، **یکپارچه‌سازی این روش با بلاکچین** است. با ثبت هش تصاویر و داده‌های جاسازی شده در یک زنجیره بلاکچین، امکان تغییر یا دستکاری اطلاعات در هر مرحله از چرخه رسیدگی به خسارت به‌طور کامل از بین می‌رود و شفافیت بیشتری ایجاد می‌شود.

در نهایت، بررسی **کارایی سیستم در مقیاس عملیاتی** و ارزیابی عملکرد آن در محیط‌های واقعی اپلیکیشن‌های بیمه‌ای، شامل زمان پردازش روی دستگاه‌های موبایل و هزینه محاسباتی

کنترل شده ثابت کرد که ترکیب استگنوگرافی مبتنی بر GAN با رمزنگاری AES-GCM، علاوه بر حفظ کیفیت تصویری، امنیت و قابلیت تشخیص دستکاری فراداده را نیز افزایش می‌دهد.

به‌طور خلاصه، روش پیشنهادی ترکیبی از **کیفیت بالا، امنیت داده‌ها، و قابلیت عملیاتی** را فراهم می‌کند و می‌تواند در مقیاس بزرگ‌تر، نقش مؤثری در کاهش تقلب‌های بیمه‌ای ایفا کند. پیشنهاد می‌شود در مراحل بعدی، این سامانه با داده‌های متنوع‌تر و در بسترهای واقعی شرکت‌های بیمه آزمایش شود و همچنین به سایر مازول‌های هوش مصنوعی مانند تحلیل محتوای تصویر یا مکان‌یابی دقیق‌تر متصل گردد. چنین توسعه‌ای می‌تواند مسیر تحول دیجیتال در صنعت بیمه را تسریع کند و الگویی برای سایر خدمات دیجیتال در کشور فراهم آورد.

جمع‌بندی و پیشنهادها

محدودیت‌ها و چالش‌های عملی روش پیشنهادی

با وجود نتایج مطلوب و برتری روش ترکیبی مبتنی بر GAN و AES-GCM نسبت به روش‌های پایه، این پژوهش نیز مانند هر سامانه مبتنی بر یادگیری عمیق، با مجموعه‌ای از محدودیت‌ها و چالش‌های عملی مواجه است که باید در توسعه‌های آتی مورد توجه قرار گیرد.

نخست، **وابستگی مدل به حجم داده‌های آموزشی** از چالش‌های اصلی در استفاده از GAN است. اگرچه در این پژوهش از ۱۰۰ تصویر واقعی خسارت بیمه‌ای برای آموزش و ارزیابی استفاده شد، اما شبکه‌های مولد متخاصم معمولاً برای رسیدن به پایداری و جلوگیری از پدیده فروپاشی مدل^۱ به داده‌های حجیم‌تر و متنوع‌تر نیاز دارند. در حجم‌های داده کوچک، احتمال همگرایی ناپایدار یا تولید الگوهای تکراری در بخش مولد افزایش می‌یابد، که ممکن است در کاربردهای واقعی به کاهش دقت یا نوسان در کیفیت جاسازی منجر شود.

دوم، **حساسیت مدل به تنظیمات هایپرپارامترها و معماری** یکی دیگر از چالش‌هاست. پارامترهایی مانند نرخ یادگیری، اندازه batch، و وزن‌های تابع هزینه در پایداری و کیفیت خروجی تأثیر مستقیم دارند. تغییرات جزئی در این مقادیر می‌تواند موجب افت PSNR یا افزایش BER شود. به همین دلیل، تنظیم دقیق و مرحله‌به‌مرحله پارامترها برای دستیابی به عملکرد پایدار ضروری است.

سوم، هزینه محاسباتی و زمانی روش GAN نسبت به LSB بیشتر است. درحالی‌که روش LSB می‌تواند در کمتر از ۰.۱۲ ثانیه داده را در هر تصویر جاسازی کند، زمان متوسط برای مدل GAN حدود ۰.۳۵ ثانیه و برای روش ترکیبی GAN+AES حدود ۰.۳۷ ثانیه برآورد شد. همچنین آموزش اولیه شبکه حدود ۴۰ دقیقه بر روی کارت گرافیکی NVIDIA RTX 3060 به طول انجامید. هرچند این زمان در محیط عملیاتی (پس از آموزش) به‌صورت برخط کمتر می‌شود، اما برای استقرار در اپلیکیشن‌های بیمه‌ای، بهینه‌سازی کد و استفاده از GPU

کپی‌رایت نویسنده(ها): © 2026 این مقاله تحت مجوز بین‌المللی Creative Commons Attribution 4.0 و CC BY اجازه استفاده، اشتراک‌گذاری، اقتباس، توزیع و تکثیر را در هر رسانه یا قالبی مشروط بر درج نحوه دقیق دسترسی به مجوز CC و منوط به ذکر تغییرات احتمالی در مقاله می‌داند. از این‌رو، به استناد مجوز یادشده، درج هرگونه تغییرات در تصاویر، منابع و ارجاعات یا دیگر مطالب از اشخاص ثالث در این مقاله باید در این مجوز گنجانده شود، مگر اینکه در راستای اعتبار مقاله به اشکال دیگری مشخص شده باشد. در صورت درج نکردن مطالب یادشده یا استفاده‌ای فراتر از مجوز فوق، نویسنده ملزم به دریافت مجوز حق نسخه‌برداری از شخص ثالث است.

به‌منظور مشاهده مجوز بین‌المللی Creative Commons Attribution 4.0 به نشانی زیر مراجعه شود:

<http://creativecommons.org/licenses/by/4.0>

یادداشت ناشر

ناشر نشریه پژوهشنامه بیمه با توجه به مرزهای حقوقی در نقشه‌های منتشرشده بی‌طرف باقی می‌ماند.

منابع

- آبرومند، ن.، مؤمن بهادر، ف.، و باقری نجیوانلو، ب. (۱۳۹۴). بهبود استگانوگرافی به روش *LSB* بر پایه الگوریتم ژنتیک چندهدفه [مقاله کنفرانسی]. دومین کنفرانس بین‌المللی و سومین کنفرانس ملی کاربرد فناوری‌های نوین در علوم مهندسی. <https://civilica.com/doc/501720>
- احمدلو، ی.، پورابراهیمی، ع.، تنها، ج.، و رجب‌زاده، ع. (۱۴۰۰). ارائه یک مدل ترکیبی برای تشخیص ادعاهای مشکوک خسارت در بیمه کشاورزی. پژوهشنامه بیمه، ۱۲(۱)، ۷۱-۳۳. <https://doi.org/10.22056/ijir.2023.01.06>
- افسرده، ح.، و دادور، م. (۱۳۹۴). استگانوگرافی به همراه کریپتوگرافی و فشرده‌سازی بر پایه شبکه‌های عصبی *MLP* [مقاله کنفرانسی]. کنفرانس بین‌المللی یافته‌های نوین پژوهشی در مهندسی برق و علوم کامپیوتر. <https://civilica.com/doc/404648>
- حسنی اژدری، س. م.، محمودزاده، آ.، خویش، م.، و رضایی، م. ع. (۱۴۰۰). نهان‌نگاری تصویر دیجیتال با استفاده از ترکیب الگوریتم ژنتیک و روش طیف گسترده در حوزه تبدیل کسینوسی گسسته فصلنامه علوم و فناوری دریا، ۲۵(۹۷)، ۱۴-۳۳. <https://doi.org/10.22034/ijmst.2021.39044>
- سخنور، ع.، و بابایی، پ. (۱۳۹۵). استفاده از شبکه عصبی فازی، الگوریتم ژنتیک و بیت کم‌اهمیت در استگانوگرافی تصویر [مقاله کنفرانسی]. مجموعه مقالات نخستین همایش ملی فناوری در مهندسی کاربردی، دانشگاه پژوهشگران جوان و نخبگان دانشگاه آزاد اسلامی (NC-TAE2016) واحد تهران غرب، ۲۱ بهمن ماه ۱۳۹۵. <https://www.sid.ir/files/server/sf/8071395h0134.pdf>
- عبدالهی، ب.، هراتی، ا.، و طاهری‌نیا، ا. (۱۴۰۲). مروری بر روش‌های نهان‌نگاری تصویر سازگار با محتوا. فصلنامه پردازش سیگنال و داده‌ها، ۲۰(۳)، ۱۸۲-۱۴۱. <https://www.noormags.ir/view/fa/2214150>

سرورهای بیمه، از جمله موضوعات مهمی است که می‌تواند زمینه‌ساز پیاده‌سازی صنعتی این پژوهش در آینده باشد.

مشارکت نویسندگان

مونا پرستش: مفهوم و طرح مقاله، جمع‌آوری و اخذ داده‌ها، تحلیل و تفسیر داده‌ها، تحلیل آماری، پیش‌نویس و اصلاح مقاله.

تشکر و قدردانی

بدین‌وسیله از شرکت معظم بیمه البرز، که به‌واسطه فعالیت در آن شرکت، از روند اعلام خسارت و دیگر فعالیت‌های بیمه‌ای مطلع هستم، و در شکل‌گیری این طرح و نگارش این مقاله مرا یاری کرد، کمال تشکر را دارم.

تعارض منافع

نویسندگان اعلام می‌دارند که در خصوص انتشار این مقاله تضاد منافع وجود ندارد. علاوه‌براین، موضوعات اخلاقی، از جمله سرقت ادبی، رضایت آگاهانه، سوءرفتار، جعل داده‌ها، انتشار و ارسال مجدد و مکرر و همچنین، سیاست مجله در قبال استفاده از هوش مصنوعی از سوی نویسندگان رعایت شده است.

دسترسی آزاد

- نوروزی، ع. (۱۴۰۴). چارچوب نهان‌نگاری استگانوگرافی قوی برای ایمنی‌سازی داده‌های حساس از راه دور با استفاده از شبکه عصبی کانولوشنال. مجله نخبگان علوم و مهندسی، ۱۰(۲)، ۱۴۲-۱۴۹. <https://elitesjournal.com/fa/page.php?rid=963>
- Abdollahi, B., Harati, A., & Taherinia, A. (2023). A review of content adaptive image steganography methods. *Signal and Data Processing Quarterly*, 20(3), 141–182. <https://www.noormags.ir/view/fa/2214150> [In Persian].
- Aberoomand, N., Momen Bahador, F., & Bagheri Najjavanloo, B. (2015). *Improving LSB-based steganography using multi-objective genetic algorithm* [Conference Presentation]. The Second International Conference and the Third National Conference on the Application of New Technologies in Engineering Sciences. <https://civilica.com/doc/501720/> [In Persian].
- Afsoorde, H., & Dadvar, M. (2015). *Image steganography combined with cryptography and compression based on MLP neural networks* [Conference Presentation]. International Conference on Novel Research Findings in Electrical and Computer Engineering. <https://civilica.com/doc/404648> [In Persian].
- Ahmadlou, Y., Pourebrahimi, A., Tanha, J., & Rajabzadeh, A. (2023). Presenting a hybrid model for identifying claims of suspicious damages in agricultural insurance. *Iranian Journal of Insurance Research*, 12(1), 63-78. <https://>

- doi.org/10.22056/ijir.2023.01.06 [In Persian].
- Dalal, M., & Juneja, M. (2021). Steganography and steganalysis (in digital forensics): A cybersecurity guide. *Multimedia Tools and Applications*, 80(4), 5723–5771. <https://doi.org/10.1007/s11042-020-09929-9>
- Ghoul, S., Sulaiman, R., & Shukur, Z. (2023). A review on security techniques in image steganography. *International Journal of Advanced Computer Science and Applications*, 14(6), 361–385. <https://doi.org/10.14569/IJACSA.2023.0140640>
- Gomathi, S., & Radhika, C. (2025). A secure messaging application using steganography and AES encryption: A dual-layer secure messaging system. *The Scientific Temper*, 16(2), 3803–3811. <https://doi.org/10.58414/SCIENTIFICTEMPER.2025.16.2.12>
- Goodfellow, I., Pouget-Abadie, J., Mirza, M., Xu, B., Warde-Farley, D., Ozair, S., Courville, A., & Bengio, Y. (2014). *Generative Adversarial Nets* [Conference Presentation]. NIPS'14: Proceedings of the 28th International Conference on Neural Information Processing Systems (Vol. 2, pp. 2672 - 2680). <https://dl.acm.org/doi/10.5555/2969033.2969125>
- Hassani Azhdari, M., Mahmoodzadeh, A., Khishe, M., & Rezaii, M. A. (2021). Digital image watermarking using the combination of genetic algorithm and spread spectrum method in the field of discrete cosine transform. *Iranian Journal of Marine Science and Technology*, 25(1), 14-33. <https://doi.org/10.22034/ijmst.2021.39044> [In Persian].
- Jaybhaye, S., Doshi, H., Dudul, A., Gaikwad, N., & Hole, P. (2023). Secure message transmission: Integrating AES encryption and LSB substitution steganography in a Django-based system. *International Journal for Research in Applied Science and Engineering Technology*, 11(12), 198–205. <https://doi.org/10.22214/ijraset.2023.57282>
- Kadhim, I. J., Premaratne, P., Vial, P. J., & Halloran, B. (2019). Comprehensive survey of image steganography: Techniques, evaluations, and trends in future research. *Neurocomputing*, 335, 299–326. <https://doi.org/10.1016/j.neucom.2018.06.075>
- Khan, N., Haan, R., Boktor, G., McComas, M., & Daneshi, R. (2020). Steganography GAN: Cracking steganography with cycle generative adversarial networks. *arXiv:2006.04008*. <https://doi.org/10.48550/arXiv.2006.04008>
- Kim, H., & Seo, H. (2025). Optimizing AES-GCM on 32-Bit ARM Cortex-M4 Microcontrollers: Fixslicing and FACE-based approach. *ACM Transaction on Embedded Computing Systems*, 24(6), 1-24. <https://doi.org/10.1145/3766074>
- Kim, K., Choi, S., Kwon, H., Kim, H., Liu, Z., & Seo, H. (2020). PAGE—practical AES GCM encryption for low end microcontrollers. *Applied Sciences*, 10(9), 3131. <https://doi.org/10.3390/app10093131>
- Klemm, R., & Chen, B. (2024). Hiding sensitive information using PDF steganography. *arXiv:2405.00865*. <https://doi.org/10.48550/arXiv.2405.00865>
- Knöchel, M., & Karius, S. (2024). *Text steganography methods and their influence in malware: A comprehensive overview and evaluation* [Conference presentation]. Proceedings of the 2024 ACM Workshop on Information Hiding and Multimedia Security. <https://doi.org/10.1145/3658664.3659637>
- Kumari, G., & Indra, G. (2024). *Video steganography with deep learning* [Conference presentation]. International Conference on Innovative Computing & Communication -ICICC 2024. <https://ssrn.com/abstract=4814247>
- Mandal, P. C., Mukherjee, I., Paul, G., & Chatterji, B. N. (2022). Digital image steganography: A literature survey. *Information Sciences*, 609, 1451–1488. <https://doi.org/10.1016/j.ins.2022.07.120>
- Nezami, Z. I., Ali, H., Asif, M., Aljuaid, H., Hamid, I., & Ali, Z. (2022). An efficient and secure technique for image steganography using a hash function. *PeerJ Computer Science*, 8, e1157, 1-18. <https://doi.org/10.7717/peerj-cs.1157>
- Norouzi, E. (2025). A robust steganography framework for securing sensitive remote data using convolutional neural networks. *Journal of Science and Engineering Elites*, 10(2), 142-149. <https://elitesjournal.com/fa/page.php?rid=963> [In Persian].
- Paul, T., Ghosh, S., & Majumder, A. (2022). A study and review on image steganography. In S. Smys, R. Bestak, R. Palanisamy, & I. Kotuliak (Eds.), *Computer Networks and Inventive Communication Technologies. Lecture Notes on Data Engineering and Communications Technologies* (Vol. 75, pp. 523-531). Springer. https://doi.org/10.1007/978-981-16-3728-5_40
- Rahman, S., Masood, F., Ullah Khan, W., Ullah, N., Qudus Khan, F., Tsaramiris, G., Jan, S., & Ashraf, M. (2020). A novel approach of image steganography for secure communication based on LSB substitution technique. *Computers, Materials & Continua*, 64(1), 31–61. <https://doi.org/10.32604/cmc.2020.09186>
- Singh, S., Singh, R., & Siddiqui, T. J. (2016). Singular value decomposition based image steganography using integer wavelet transform. In S. M. Thampi, S. Bandyopadhyay, S. Krishnan, K.-C. Li, S. Mosin, & M. Ma (Eds.), *Advances in Signal Processing and Intelligent Recognition Systems*. Advances in Intelligent Systems and Computing (Vol. 425, pp. 593-601). Springer. https://doi.org/10.1007/978-3-319-28658-7_50
- Sokhanvar, A., & Babaei, P. (2016). *Application of fuzzy neural networks, genetic algorithms, and least significant bit in image steganography* [Conference Presentation]. Proceedings of the First National Conference on Technology in Applied Engineering (NCTAE2016), Young Researchers and Elites Club, Islamic Azad University, Tehran West Branch. <https://www.sid.ir/files/server/sf/8071395h0134.pdf> [In Persian].
- Subramanian, N., Elharrouss, O., Al Maadeed, S., & Bouridane, A. (2021). Image steganography: A review of the recent advances. *IEEE Access*, 9, 23409-23423. <https://doi.org/10.1109/ACCESS.2021.3053998>
- Tushara, M., & Navas, K. A. (2016). Image steganography using discrete wavelet transform— A review. *International Journal of Innovative Research in Electrical, Electronics, Instrumentation and Control Engineering*, 3(1s), 207–212. <https://ijireeice.com/wp-content/uploads/2016/07/nCORETech-38.pdf>
- Venkat Sudheer, P., Harini, M., & Jayachandra, G. (2025). GAN powered image steganography: Combining NLP and generative adversarial networks for text and voice encryption. *International Journal of Science and*

Advanced Technology (IJSAT), 16(2), 1-12. <https://doi.org/g9r8fj>
Viaene, S., & Dedene, G. (2004). Insurance fraud: Issues and challenges. *The Geneva Papers on Risk and Insurance - Issues and Practice*, 29(2), 313–333.

<https://doi.org/10.1111/j.1468-0440.2004.00290.x>
Zhang, K. A., Cuesta Infante, A., Xu, L., & Veeramachaneni, K. (2019). SteganoGAN: High-capacity image steganography with GANs. *arXiv:1901.03892*. <https://doi.org/10.48550/arXiv.1901.03892>

AUTHOR(S) BIOSKETCHES	معرفی نویسندگان
مونا پرستش (Mona Parastesh)، دانشجوی دکتری، گروه مهندسی کامپیوتر، دانشکده فنی و مهندسی، دانشگاه آزاد اسلامی، واحد تهران جنوب، تهران، ایران. - Email: St_m_parastesh@azad.ac.ir - ORCID: 0009-0008-2851-7702	

HOW TO CITE THIS ARTICLE	
Parastesh, M. (2026). A novel approach to reducing fraud in insurance processes using GAN-Based steganography and AES-GCM encryption. <i>Iranian Journal of Insurance Research.</i> , 15(2), 119-140.	
DOI: 10.22056/ijir.2026.02.02 URL: https://ijir.irc.ac.ir/article_160371.html	

پژوهشگاه علوم انسانی و مطالعات فرهنگی
پرتال جامع علوم انسانی