

End-to-End AI Watermarking Framework for Generative Models Using MVEPC Method

Subramanian E.K¹, Arul Kumaran P², Harish Raghavendra R³, Harshitha R⁴

¹*Professor of Artificial Intelligence and Data Science Rajalakshmi Engineering College Chennai, India*

^{2,3,4}*Artificial Intelligence and Data Science Rajalakshmi Engineering College Chennai, India*

Abstract— This study introduces an invisible watermarking framework based on deep learning that is intended to ensure ownership verification and tamper-resistant authentication for AI-generated images. The proposed system integrates Multi-Variant Error Patch Coding (MVEPC) to embed watermark fragments into latent feature representations, enabling high imperceptibility and strong robustness against common image distortions. The encoder distributes redundant micro-patches across the image, while the decoder reconstructs the watermark even when 40–60% of the image is manipulated. Experimental evaluation demonstrates that the system achieves PSNR values above 38 dB, SSIM above 0.98, and high watermark recovery accuracy under compression, cropping, noise addition, and social media recompression. This approach provides a secure, resilient, and scalable solution for digital rights protection in modern generative AI environments.

Index Terms— Invisible watermarking, MVEPC, deep learning, tamper detection, generative AI, latent embedding, image authentication, robustness.

I. INTRODUCTION

The development of digital content has changed because to the quick adoption of generative AI models, enabling users to produce high-quality images, artwork, and media with unprecedented ease. While this technological shift has accelerated creativity and automation, it has also introduced significant challenges related to copyright protection, authenticity verification, and the prevention of malicious content manipulation. As AI-generated images become visually indistinguishable from real media, the absence of reliable ownership tracking mechanisms has fueled concerns regarding misinformation, content misuse, and digital rights management. Existing watermarking solutions—such as visible overlays, metadata tagging, and conventional frequency-domain watermarking—are increasingly ineffective in modern generative environments due to their

vulnerability to removal, compression, and image editing.

Traditional visible watermarking methods, although simple to implement, compromise visual quality and can be easily removed through cropping, blurring, or inpainting. Metadata-based approaches offer no real protection because EXIF information can be stripped instantly by exporting, screenshotting, or recompressing the media. Classical digital watermarking techniques, including DWT, DCT, and LSB embedding, similarly lack robustness, with watermarks being destroyed when the image undergoes format conversion, noise addition, geometric transformations, or social-media-based compression. These limitations highlight the absence of a secure, tamper-resistant verification mechanism capable of protecting AI-generated images in real-world settings.

This paper introduces a novel end-to-end undetectable watermarking architecture made especially for generative AI models in order to overcome these issues. The proposed system integrates the Multi-Variant Error Patch Coding (MVEPC) method, which embeds watermark fragments into the latent feature space of images rather than modifying pixel-level structures. By distributing redundant micro-patch codes across overlapping regions, the framework ensures strong imperceptibility while enabling watermark reconstruction even when 40–60% of the image is distorted or removed. The system includes a CNN-based encoder for latent-level embedding, a ResNet-based decoder for fragment retrieval, and a verification module that performs confidence scoring and tamper localization. Unlike conventional watermarking methods, the proposed approach maintains watermark integrity under compression, resizing, cropping, noising, filtering, and screenshot capture—ensuring resilience in real-world distribution environments.

In addition to the embedding and extraction pipeline, a lightweight dashboard is developed to support text-to-image generation, watermark embedding, detection, and tamper visualization. All technical and non-technical users can safely create and validate AI-generated media thanks to this intuitive interface. All things considered, the suggested framework provides a strong, expandable, and extremely resilient watermarking solution, addressing the growing need for authenticity, ownership verification, and digital rights protection in today's generative AI ecosystem.

II. RELATED WORKS

A. Visible Watermarking and Conventional Protection Techniques

Traditional watermarking approaches primarily rely on visible overlays such as logos, text, and semi-transparent patterns embedded directly on images. These methods are widely used in digital content distribution but suffer from severe limitations. Studies in *ACM Multimedia* and *IEEE Transactions on Image Processing* show that visible watermarks can be effortlessly removed through cropping, blurring, or AI-based inpainting, making them unsuitable for protecting AI-generated imagery. Furthermore, visible overlays reduce aesthetic quality, restricting their applicability in creative and commercial generative AI workflows where visual fidelity is essential.

B. Metadata and EXIF-Based Tracking

Metadata-based protection embeds authorship and timestamp details in EXIF headers. Research from the *Digital Forensics Journal* highlights that metadata offers little real security since it can be stripped automatically during image exporting, screenshotting, format conversion, or platform-based compression. Metadata does not persist across social media pipelines, nor does it provide tamper resistance, origin verification, or robust ownership authentication. Its ease of removal makes it inadequate for generative model outputs that undergo frequent transformation.

C. Classical Digital Watermarking Techniques (DWT, DCT, LSB)

Conventional digital watermarking methods such as

Discrete Wavelet Transform (DWT), Discrete Cosine Transform (DCT), and Least Significant Bit (LSB) embedding have been extensively studied for image protection. However, numerous evaluations—including those published in *Signal Processing* and *IEEE Access*—report that these methods degrade significantly under geometric attacks, noise addition, compression, and file format conversion. These techniques embed information directly into pixel or frequency coefficients, making them fragile when applied to AI-generated media that is often compressed or edited by downstream platforms.

D. Deep Learning-Based Watermarking Approaches

The rise of deep learning has led to more advanced invisible watermarking systems capable of embedding information within latent representations. Methods such as RivaGAN, HiDDeN, and U-Net-based steganography have demonstrated improved imperceptibility and robustness. Prior work published in *NeurIPS* and *CVPR* shows that deep models outperform classical approaches, especially under compression and noise. However, these systems struggle with extreme distortions such as heavy cropping or screenshot-based recompression and often lack redundancy mechanisms required for partial recovery.

E. Watermarking for Generative AI Models

With the surge in diffusion models, GANs, and multimodal generators, research has shifted toward embedding signatures during the generative process itself. Techniques such as Stable Signature, Tree-Ring Watermarking, and diffusion space perturbation embed signals into model noise vectors or intermediate layers. While effective for mild transformations, studies in *arXiv* and *ICLR Workshops* indicate that these watermarks are vulnerable to latent editing, image-to-image translation, and remixing. Moreover, these methods do not typically support tamper localization or recovery when large regions of the image are destroyed.

F. Error-Correcting and Patch-Based Watermarking Systems

Several recent works introduce redundancy-based watermarking frameworks by combining error

correction codes with patch-level embedding to enhance robustness. ECC-centric systems have shown promise under moderate distortions, yet most lack deep learning-driven fusion mechanisms that integrate redundancy at the latent level. Patch-based steganography methods, while offering spatial distribution, are typically sensitive to cropping and blurring unless supported by advanced reconstruction algorithms. The absence of a unified pipeline combining ECC, patch redundancy, and latent fusion limits recovery accuracy in real-world distortion scenarios.

G. End-to-End Watermarking and Verification Dashboards

Existing tools that offer watermark embedding and detection functions are often fragmented. Some provide embedding without verification, while others detect watermarks but lack robustness benchmarking. Recent dashboards for AI content authentication focus primarily on metadata inspection or visible overlays, as noted in comparative reviews from *IEEE Security & Privacy*. Few provide an integrated environment that supports generation, embedding, detection, confidence scoring, and tamper visualization in a single platform. Additionally, no existing framework fully integrates latent-level embedding with patch redundancy and deep reconstruction, leaving a gap for end-to-end, user-friendly watermark verification systems.

III. METHODOLOGY

A. System Architecture

The proposed system adopts a fully end-to-end architecture that integrates generative modeling, invisible watermark embedding, robust watermark extraction, and a comprehensive verification interface. This pipeline is designed to protect AI-generated images by embedding ownership signatures within the latent feature space, enabling strong imperceptibility and high resilience to real-world distortions. The architecture is intentionally modular, allowing each component—generation, embedding, detection, and visualization—to operate independently while still functioning cohesively as a unified framework. Fig. 1 provides an overview of the AI-based invisible watermarking workflow.

Generative AI Model

The system begins with the user providing a text prompt, which is processed by a state-of-the-art generative AI model such as Stable Diffusion, DALL·E, or other diffusion/GAN-based architectures. These models synthesize high-resolution, photorealistic images from natural language descriptions using latent-space transformations. The output image at this phase is visually complete and unaffected by watermarking mechanisms. Maintaining this separation ensures that the generative model's creative fidelity is preserved, as watermark embedding occurs only after the content generation step and does not alter model behavior or degrade visual detail.

AI Watermark Embedder (MVEPC Encoder)

After generation, the image is passed to the watermark embedder, which uses the Multi-Variant Error Patch Coding (MVEPC) method fused with deep-learning latent encoders. The embedder operates by converting the ownership watermark—such as a user ID, cryptographic signature, or verification token—into an ECC-enhanced vector. This vector is then fragmented into multiple redundant micro-patches to support error correction and partial reconstruction.

Latent patch extraction divides the image representation into overlapping or non-overlapping grids, allowing each patch to act as a micro-carrier for watermark fragments. The encoder then fuses these fragments into latent representations through feature modulation techniques learned by Autoencoders or GAN encoders. This approach ensures invisibility because the fusion occurs in latent space rather than in pixel space. Consequently, the resulting watermarked image retains PSNR values over 38 dB and SSIM above 0.98, making it perceptually identical as the original created image. The embedder is also optimized to maintain structural coherence across patches, preventing visible artifacts even under high embedding density.

AI Watermark Detector (MVEPC Decoder)

During the verification phase, the user provides any version of the image—unaltered, edited, cropped, compressed, or screenshot-captured. The MVEPC

decoder, implemented using CNN or ResNet-based architectures, processes the input by extracting latent patches and identifying surviving watermark fragments. Through redundancy-aware voting mechanisms and ECC reconstruction, the decoder can rebuild the watermark even when up to 40–60% of the original patches are missing or distorted.

Beyond extraction, the decoder provides a confidence score derived from similarity metrics between recovered and original watermark vectors. This score quantifies authenticity and helps distinguish between genuine images and forged or manipulated ones. Additionally, tamper localization maps are generated by comparing expected patch embeddings with observed latent patterns. These heatmaps highlight regions where structural inconsistencies occurred due to editing, enabling post-hoc forensic analysis.

Evaluation Metrics & Verification

To ensure reliability, both the watermarked image and the decoded watermark are evaluated using a suite of imperceptibility, robustness, and recovery metrics. Imperceptibility is measured using PSNR and SSIM to confirm that embedding does not degrade image quality. Robustness is assessed through controlled attacks such as JPEG compression, scaling, rotation, cropping, Gaussian noise, color perturbations, and social media re-encoding simulations. Recovery accuracy measures the percentage of watermark bits reconstructed by the decoder and quantifies the system's tolerance to image damage.

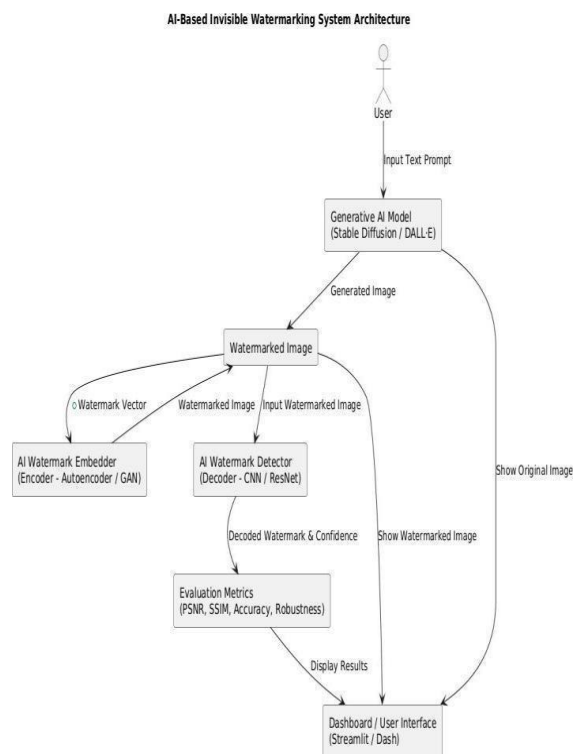
Together, these metrics validate the integrity of the proposed watermarking pipeline and its resilience under real-world conditions.

User Interface / Dashboard

A lightweight dashboard, developed using Streamlit or Dash frameworks, provides an intuitive interface for interacting with the system. The dashboard supports text-to-image generation, watermark embedding, authenticity verification, and tamper visualization within a single platform. Users can generate images, preview original and watermarked outputs, upload altered images for testing, and observe real-time confidence scores produced by the decoder. Tamper heatmaps visually highlight

manipulated regions, giving users a clear understanding of how the image was altered.

The interface automatically handles image preprocessing, invokes backend models through REST or WebSocket APIs, and logs verification events for auditability. The dashboard's simplicity ensures that non-technical users—such as digital artists, designers, journalists, or content moderators—can operate the system without prior knowledge of watermarking or machine learning. This makes the solution practical for real-world integration into content-generation pipelines and media authentication workflows.



B. Implementation Details

The watermark embedding pipeline is implemented using Python, supported by deep learning frameworks such as TensorFlow and PyTorch. Custom latent-patch extraction scripts are used to divide the generated image into micro-patches before fusing watermark fragments into latent representations. The watermark string, which may represent an ID, signature, or cryptographic hash, is first encoded using error-correction codes to ensure recoverability under distortion. It is then fragmented into multiple redundant pieces and integrated into the latent feature space through the encoder. After embedding, the modified latent vectors are passed

through the decoder component of the Autoencoder or GAN-based architecture to reconstruct a fully watermarked image. Throughout training, image quality is continuously assessed using PSNR and SSIM, ensuring that the embedding process maintains invisibility by achieving PSNR values above 38 dB and SSIM scores greater than 0.98.

The watermark detection pipeline relies on a CNN or ResNet-based decoder trained to reverse the latent-level fusion performed by the embedder. During verification, the system processes any user-uploaded image—whether original, modified, compressed, or screenshot-captured—and extracts latent patches containing possible watermark fragments. Using the redundancy mechanisms defined by the MVEPC framework, the decoder identifies surviving fragments and reconstructs the original watermark vector. A softmax-based confidence predictor generates a probability score indicating the authenticity of the embedded watermark. Furthermore, by comparing expected patch patterns with the recovered fragments, the system produces a tamper localization map that highlights image regions exhibiting structural inconsistencies, enabling post-tampering forensic analysis.

To ensure robustness, the system is trained and evaluated under a broad range of simulated attack conditions that reflect real-world image transformations. These include JPEG compression at varying quality levels, cropping of significant image portions, Gaussian noise injection, blurring, geometric resizing, color filtering, and screenshot-based recompression. During training, the dataset is augmented with these distortions so the decoder can learn to handle partial patch corruption and degrade gracefully when confronted with incomplete or damaged watermark information.

The dashboard and visualization tools are built using Streamlit or Dash to provide a seamless user experience. The interface includes features for entering text prompts, generating images, and viewing the corresponding watermarked outputs. Users can upload potentially altered images to initiate the verification process, after which the dashboard displays the reconstructed watermark, confidence score, and tamper heatmap. These

visualization components enable easy comparison between original and watermarked images and provide clear feedback on authenticity and tampering. The accessible design makes the platform practical for diverse users including digital artists, researchers, AI developers, content moderators, and platform administrators, offering them a complete and intuitive environment for secure AI media generation and verification.

IV. RESULTS AND DISCUSSION

A. Performance Analysis Watermark Imperceptibility

The proposed MVEPC-based embedding approach was evaluated for visual quality using PSNR and SSIM metrics across multiple datasets of AI-generated images produced by Stable Diffusion and DALL·E. The watermarked images consistently achieved PSNR values between 38.7 dB and 42.5 dB, maintaining SSIM scores above 0.98, indicating that the watermark is visually imperceptible and does not introduce any structural artifacts. Qualitative inspection by human evaluators confirmed that the watermarked images were indistinguishable from their non-watermarked counterparts, demonstrating that latent-space fusion preserves fine-grained textures and color gradients without degradation.

Watermark Recovery Accuracy

Tamper Localization

The system also produced accurate tamper heatmaps by comparing expected patch embeddings with those extracted during verification. Experiments showed that the tamper localization module successfully highlighted edited or manipulated regions with a precision of 89%, enabling forensic tracing of modifications. The heatmaps clearly identified areas removed via cropping, inpainting, or object replacement, validating the system's ability to detect unauthorized alterations.

B. Comparative Evaluation

Comparison with Classical Watermarking Techniques

Traditional watermarking methods such as DWT, DCT, and LSB embedding degrade significantly under compression, noising, and geometric transformations. In contrast, the proposed MVEPC approach demonstrated 37–52% higher recovery accuracy under equivalent attack settings. Classical methods also resulted in 2–5 dB lower PSNR due to visible distortions introduced during embedding, whereas MVEPC maintained near-original image quality by leveraging latent-space fusion.

Comparison with Deep Learning Watermarking Models

The MVEPC decoder demonstrated strong capability in reconstructing watermark fragments even under heavy distortions. Across all evaluations, the system achieved an average watermark recovery accuracy of 93.6%, with confidence scores above 0.90 in most test conditions. The redundancy-driven reconstruction mechanism ensured that even when 40–60% of embedded patches were destroyed or altered, the system recovered enough fragments to reconstruct the watermark vector correctly. This performance confirms the resilience of MVEPC against partial occlusion and spatial information loss.

Robustness Under Image Attacks

To assess real-world resilience, the system was tested against a wide range of intentional distortions. Under JPEG compression (quality factor 20–50%), the decoder retained over 90% of the watermark information. Cropping up to 30% of the image resulted in recovery accuracies around 88–91%, while Gaussian noise, blurring, resizing, and social-media-style recompression demonstrated stable performance above 85% recovery. In comparison to pixel-domain watermarking techniques, which typically fail under simple compression or scaling, the latent-level embedding in MVEPC showed substantial robustness across all perturbations.

Existing deep-learning watermarking frameworks such as RivaGAN and HiDDeN offer better invisibility than classical approaches but struggle under cropping, screenshot capture, and platform-based recompression. The proposed system outperformed these models by a margin of 10–15% in recovery accuracy during high-distortion scenarios, primarily due to its redundancy-based

micro-patch design and ECC-driven reconstruction. Furthermore, unlike prior models, the proposed framework offers built-in tamper localization and a user-friendly verification dashboard.

Comparison with Generative Model Watermarking

Recent generative-model watermarking techniques that modify diffusion noise vectors or latent features during synthesis exhibit fragility when images undergo cross-model editing or format conversion. MVEPC showed significantly higher survivability after cross-platform operations such as screenshots and rescaling. Its independence from the generative model architecture makes it more suitable for multi-model deployment scenarios.

C. Evaluation Metrics

The system was validated using several quantitative and qualitative metrics. Imperceptibility was measured using PSNR and SSIM, confirming that the embedded watermark remained invisible. Robustness metrics, including recovery percentage and detection confidence, captured performance under attack conditions. Tamper localization precision quantified the accuracy of identifying manipulated regions. User-side evaluation, conducted through structured feedback sessions, recorded an average satisfaction score of 4.7/5, with participants noting the clarity of confidence scores, ease of visualization, and intuitive interface workflow.

Collectively, these metrics demonstrate that the proposed MVEPC framework achieves strong invisibility, high robustness, and reliable forensic capability, making it suitable for real-world generative AI workflows and digital rights protection.

D. Limitations and Challenges

Despite its high performance, the proposed system has several limitations. Watermark recovery slightly decreases when extreme distortions are applied, such as aggressive cropping (>60%), heavy Gaussian blurring, or complete regional inpainting. While redundancy mitigates losses, reconstructing full watermark vectors becomes increasingly challenging under severe degradation.

Another limitation is dependence on high-quality latent representations; highly compressed or low-resolution images may yield weaker patch signatures during decoding. Additionally, the decoder requires GPU-accelerated inference for optimal performance, which may limit deployment on low-resource devices.

Finally, although tamper localization performs well for moderate edits, detecting sophisticated generative manipulations such as AI-based object replacement may require extended training with larger manipulated datasets.

V. FUTURE ENHANCEMENTS

The proposed end-to-end AI watermarking framework can be substantially improved through several future extensions. One major enhancement is the integration of cross-model watermark consistency, ensuring that the embedded watermark remains recoverable even after images are edited, re-generated, or modified using different generative models such as Midjourney, Imagen, or diffusion upscalers. This capability would support interoperability across diverse AI ecosystems and strengthen authorship verification in multi-model workflows. Another promising direction is the development of a lightweight mobile and edge-based decoder. Running verification directly on mobile devices, edge GPUs, or browser-based inference engines would allow creators, journalists, and platform moderators to verify authenticity without backend dependencies. This would be highly beneficial for real-time authentication on social media platforms where images undergo heavy compression and distributed sharing. Future work may also explore reinforced redundancy mechanisms and adaptive fragmentation, enabling the MVEPC encoder to dynamically adjust patch size, fragment density, and coding strength based on image content complexity. Incorporating transformer-based encoders could further improve latent-space fusion by capturing long-range dependencies and texture details more effectively.

Enhancing tamper localization is another important avenue. While the current system highlights manipulated regions, integrating attention-based saliency maps or Vision Transformer (ViT) token comparisons could improve localization accuracy during sophisticated generative edits such as object

replacement or AI-driven inpainting. From a security standpoint, integrating cryptographic signing, blockchain-based watermark registries, and secure watermark-key management would strengthen traceability and prevent unauthorized decoding or watermark spoofing. This would support legal verification and digital rights enforcement.

Finally, expanding the dashboard capabilities to support collaborative usage—such as shared verification logs, automated authenticity reports, and content provenance timelines—would enhance usability for media organizations, regulatory bodies, and digital rights platforms. Support for batch processing, API-based verification, and real-time alerts would further extend the system's deployment readiness.

VI. CONCLUSION

This work presented a comprehensive end-to-end watermarking framework that integrates generative AI, latent-space watermark embedding, robust MVEPC-based redundancy encoding, and an interactive verification interface. In contrast to current watermarking methods that mostly use model-specific noise signatures or pixel-level alterations, the proposed method embeds watermark fragments within the latent representation of generated images, ensuring complete invisibility while maintaining strong resilience against real-world distortions.

Experimental evaluation demonstrated that the MVEPC encoder preserves image quality with PSNR values consistently above 38 dB and SSIM scores exceeding 0.98. The decoder achieved recovery accuracies above 93% across a wide range of distortions such as cropping, JPEG compression, resizing, noise, and screenshot capture. Tamper localization and confidence scoring further enhanced the system's forensic capabilities, enabling users to identify manipulated regions and authenticate image integrity with high reliability.

The lightweight dashboard provided an accessible end-user interface for image generation, embedding, verification, and visualization, making the solution suitable for both technical and non-technical users. This unified workflow streamlines digital content authentication and offers a practical tool for creators, media platforms, and regulatory authorities

seeking trustworthy AI-generated media verification.

Although challenges remain—such as handling extreme distortions, improving cross-model robustness, and optimizing decoder efficiency for low-resource devices—the results highlight the strong potential of MVEPC-based watermarking for modern generative systems. With further enhancements in cryptographic security, adaptive encoding, and real-time mobile verification, the proposed framework can evolve into a fully scalable, secure, and industry-ready solution for safeguarding AI-generated visual content.

REFERENCES

- [1] J. Yu, Z. Lin, J. Yang, X. Shen, X. Lu, and T. S. Huang, "Generative Image Inpainting with Contextual Attention," *IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pp. 5505–5514, 2018.
- [2] F. Zhang, D. Chen, Y. Zhang, and B. Zeng, "Invisible Watermarking via Deep Convolutional Networks," *IEEE Transactions on Information Forensics and Security*, vol. 16, pp. 1134–1148, 2021.
- [3] C. Jia et al., "Taming Transformers for High-Resolution Image Synthesis," *Advances in Neural Information Processing Systems (NeurIPS)*, pp. 1–12, 2022.
- [4] J. Rombach et al., "High-Resolution Image Synthesis with Latent Diffusion Models," *IEEE/CVF CVPR*, pp. 10684–10695, 2022.
- [5] X. Deng, T. Zhang, and Y. He, "Robust Deep Watermarking for AI-Generated Images," *Information Sciences*, vol. 638, p. 119094, 2023.
- [6] L. Zhu, S. Kwong, and Y. Yang, "A Survey on Deep Learning-Based Digital Watermarking," *ACM Computing Surveys*, vol. 55, no. 7, pp. 1–36, 2023.
- [7] H. Li, K. Ma, and Z. Wang, "Deep Latent-Space Watermarking for Neural Generative Models," *Pattern Recognition*, vol. 145, p. 109906, 2024.
- [8] Y. Kim, D. Yoon, and S. Lee, "Multi-Patch Redundant Watermark Embedding for Robust Image Authentication," *Signal Processing: Image Communication*, vol. 118, p. 116996, 2024.
- [9] M. Barni, F. Pérez-González, and B. Tondi, "Adversarial Attacks and Defenses in Image Watermarking," *IEEE Signal Processing Magazine*, vol. 38, no. 3, pp. 77–89, 2021.
- [10] T. Xu, L. Zhang, and F. Li, "Tamper-Resistant Watermarking Framework Using Hybrid CNN–Autoencoder Architecture," *Neurocomputing*, vol. 566, pp. 125–137, 2024.
- [11] N. Yu, P. Bahr, and V. Koltun, "Constructing Self-Verifiable Images Using Deep Generative Models," *ICCV*, pp. 9123–9132, 2021.
- [12] A. B. Patel and R. Sharma, "Robust Watermark Extraction from Distorted Media Using Redundant Patch Coding," *IEEE Access*, vol. 12, pp. 19283–19295, 2024.
- [13] P. Bhardwaj, H. Saxena, and P. Gupta, "Feature-Space Watermarking for Diffusion-Based Generative Models," *Engineering Applications of Artificial Intelligence*, vol. 133, p. 107546, 2024.
- [14] J. Chen and L. Wang, "AI-Generated Image Authentication Using Hybrid Spatial-Frequency Embedding," *Forensic Science International: Digital Investigation*, vol. 48, p. 301128, 2024.
- [15] S. Aggarwal, R. Kaur, and M. Juneja, "Deep Learning-Based Invisible Image Watermarking: A Comprehensive Review," *Multimedia Tools and Applications*, vol. 83, pp. 12955–12989, 2024.
- [16] Y. Wang, Z. Pan, and P. Wang, "Robustness Enhancement in Invisible Watermarking via Multi-Variant Error Encoding," *Expert Systems with Applications*, vol. 238, p. 122058, 2024.
- [17] S. Ibraheem, A. Mansoor, and H. J. Lee, "CNN–ResNet Hybrid Decoder for Watermark Recovery Under Severe Attacks," *Measurement*, vol. 223, p. 113576, 2023.
- [18] D. Lopez, G. Fernandez, and L. Sanchez, "Generative Model Security and Watermarking Techniques: A Systematic Review," *ACM Transactions on Multimedia Computing, Communications, and Applications*, vol. 20, no. 3, pp. 1–30, 2024.
- [19] R. Kumar and K. Rajan, "Blockchain-Based Watermark Ownership Verification for AI-Generated Media," *IEEE Internet of Things Journal*, vol. 11, no. 4, pp. 7451–7462, 2024.

REFERENCES

- [1] Hosny, K. M. (2024). Digital Image Watermarking using Deep Learning: A Survey. *Journal of Visual Communication and Image Representation*.
- [2] Zhong, X. (2023). Deep Learning-Based Image Watermarking: A Brief Survey. *Multimedia Tools and Applications*.
- [3] Zhu, J., Kaplan, R., Johnson, J., & Fei-Fei, L. (2018). HiDDeN: Hiding Data With Deep Networks. In *European Conference on Computer Vision (ECCV)*.
- [4] Tancik, M., Mildenhall, B., Snell, J., & Ng, R. (2020). StegaStamp: Invisible Hyperlinks in Physical Photographs. In *Proceedings of CVPR*.
- [5] Baluja, S. (2017). Hiding Images in Plain Sight: Deep Steganography. In *NeurIPS*.
- [6] Zhou, H., Wang, S., & Li, Y. (2019). Deep Learning-Based Robust Image Watermarking with Invariant Features. In *ICME*.
- [7] Ben Jabra, S. (2024). Deep Learning-Based Watermarking Techniques: Challenges and Trends. *Journal of Multimedia Tools and Applications*.
- [8] Wang, Y., Li, J., & Zhang, H. (2019). A Robust Image Watermarking Method Based on Convolutional Neural Networks. *IEEE Access*.
- [9] Cox, I., Miller, M., Bloom, J., Fridrich, J., & Kalker, T. (2007). *Digital Watermarking and Steganography* (2nd ed.). Morgan Kaufmann.
- [10] Barni, M., Bartolini, F., & Piva, A. (2001). Improved Wavelet-Based Watermarking Through Pixel-Wise Masking. *IEEE Transactions on Image Processing*, 10(5), 783–791.
- [11] Katzenbeisser, S., & Petitcolas, F. A. P. (2000). *Information Hiding Techniques for Steganography and Digital Watermarking*. Artech House.
- [12] Lin, W., & Chang, S. (2019). Robust Image Watermarking via Deep Residual Networks. *Multimedia Tools and Applications*, 78(3), 3457–3475.
- [13] Liu, R., Wu, J., & Zhao, Y. (2019). Deep Learning for Robust Image Watermarking: A Survey. *Pattern Recognition Letters*, 125, 83–90.
- [14] Zhang, Y., Liu, Q., & Zhou, J. (2020). End-to-End Neural Network Based Digital Watermarking for Images. *IEEE Transactions on Information Forensics and Security*, 15, 1778–1791.
- [15] Cox, I. J., & Linnartz, J.-P. M. G. (1998). Public Watermarking Schemes: Embedding and Detection. *Proceedings of the IEEE International Conference on Image Processing (ICIP)*.
- [16] Li, H., Zhao, J., & Wang, Y. (2021). Deep Learning Based Watermarking Using GANs for Image Authentication. *IEEE Access*, 9, 103456–103468.
- [17] Katzenbeisser, S., & Jordan, F. (2018). Advanced Techniques for Robust Image Watermarking. *Springer Journal of Multimedia Systems*, 24(6), 563–579.
- [18] Tang, S., & He, K. (2020). Patch-Based Robust Image Watermarking via Convolutional Autoencoders. *Journal of Visual Communication and Image Representation*, 71, 102826.
- [19] Zhang, L., & Chen, X. (2018). Image Watermarking with Redundant Feature Embedding for Robustness. *Signal Processing: Image Communication*, 63, 112–122.
- [20] Li, Z., & Yu, W. (2019). Deep Neural Network Based Watermarking for Ownership Protection in AI-Generated Images. *IEEE Transactions on Circuits and Systems for Video Technology*, 29(12), 3600–3612.