

Robust Image Watermarking Against Geometric Attacks Using Surrogate Model-Based Scaling Factor Optimization and Deep Learning-Based Adaptive Embedding

Masoud. Pourmoradi¹, Hamidreza. Ghaffari^{1*}

¹ Department of Computer Engineering, Fe.C., Islamic Azad University, Ferdows, Iran

* Corresponding author email address: hamidghaffary53@iau.ac.ir

Article Info

Article type:

Original Research

How to cite this article:

Pourmoradi, M. & Ghaffari, H. (2026). Robust Image Watermarking Against Geometric Attacks Using Surrogate Model-Based Scaling Factor Optimization and Deep Learning-Based Adaptive Embedding. *Decision Science and Intelligent Systems*. 3(2), 1-37.



© 2026 the authors. Published by KMAN Publication Inc. (KMANPUB), Ontario, Canada. This is an open access article under the terms of the Creative Commons Attribution-NonCommercial 4.0 International (CC BY-NC 4.0) License.

ABSTRACT

Digital watermarking, as a fundamental approach to copyright protection, has consistently faced the inherent challenge of balancing imperceptibility and robustness, particularly against geometric attacks. This delicate trade-off is strongly influenced by two key factors: the optimal selection of the scaling factor (embedding strength) and the intelligent selection of embedding regions. This article proposes a robust hybrid watermarking framework for color images that strategically integrates intelligent image analysis with accelerated optimization and is specifically designed to withstand rotation attacks. The proposed method incorporates two principal innovations. First, a deep convolutional neural network (CNN) with a U-Net architecture is trained to generate an adaptive weight map. Independently of the embedding process, this network analyzes the content of the host image and automatically identifies regions with high visual complexity, including edges and textures. Second, a particle swarm optimization (PSO) algorithm guided by a regression tree-based surrogate model is employed to determine the optimal scaling factor. By replacing computationally expensive evaluations of the actual objective function with a rapid estimator, this approach substantially accelerates the optimization process. Finally, fully adaptive embedding is performed through the pointwise integration of the optimal scaling factor and the local weight map within the approximation subband of the discrete wavelet transform (DWT). Experimental results obtained using standard benchmark images confirm the strong performance of the proposed method. By achieving an average peak signal-to-noise ratio (PSNR) greater than 38 dB and a normalized correlation (NC) exceeding 98% following the simultaneous application of rotation and JPEG compression attacks, the proposed method demonstrates a substantial improvement over previous approaches and effectively establishes a balance between imperceptibility and robustness.

Keywords: image watermarking, geometric attacks, scaling factor optimization, surrogate model, deep learning, convolutional neural network

Extended Abstract

Introduction

The rapid expansion of digital communication technologies has dramatically increased the production, distribution, and exchange of multimedia content across social networks, electronic commerce platforms, medical databases, news websites, and other digital environments. Although this development has facilitated efficient access to digital information, it has simultaneously intensified concerns regarding unauthorized reproduction, manipulation, redistribution, copyright infringement, and ownership verification. Digital image watermarking has consequently emerged as an important complementary security mechanism through which ownership or authentication information can be imperceptibly embedded into a host image while remaining recoverable when verification is required (Mahto et al., 2021; Zermi et al., 2021). Unlike conventional cryptographic techniques, whose protection generally terminates after content is decrypted, digital watermarking maintains ownership-related information as an integral component of the digital media itself.

The effectiveness of a digital watermarking system is fundamentally determined by the trade-off among imperceptibility, robustness, and embedding capacity. Imperceptibility requires the watermarked image to remain visually indistinguishable from the original host image, whereas robustness concerns the ability to recover the embedded watermark after intentional or unintentional image processing. Capacity refers to the amount of information that can be embedded without unacceptable deterioration in image quality (Begum & Uddin, 2020). Among these requirements, robustness against geometric attacks remains particularly challenging. Rotation, scaling, and translation attacks can disrupt synchronization between the embedding and extraction procedures without necessarily destroying the watermark itself, thereby considerably reducing extraction reliability (Degadwala et al., 2020).

Transform-domain watermarking has generally demonstrated greater robustness than spatial-domain techniques because watermark information can be distributed among perceptually and structurally significant frequency coefficients. Discrete cosine transform (DCT), discrete wavelet transform (DWT), and singular value decomposition (SVD) have therefore become prominent techniques in robust watermarking systems (Abdulrahman & Ozturk, 2019; Fares et al., 2021). DWT is particularly advantageous because its multiresolution representation simultaneously provides spatial and frequency localization, allowing watermark information to be embedded into subbands with different perceptual and robustness characteristics. However, two major problems remain unresolved. The first concerns selection of the scaling factor, or embedding strength. An excessively small scaling factor preserves image quality but produces a vulnerable watermark, whereas an excessively large factor improves robustness at the expense of perceptual quality. The second concerns the conventional use of uniform embedding strength throughout an image despite substantial local differences in visual complexity.

Recent advances in deep learning provide an opportunity to address the latter limitation. Convolutional neural networks (CNNs) can learn hierarchical visual representations and distinguish smooth regions from complex edges and textures, making them suitable for identifying image areas in which watermark energy can be concealed more effectively (Zhu et al., 2021). Although end-to-end deep watermarking systems have demonstrated promising performance, their robustness against geometric transformations frequently depends on computationally intensive attack-aware training, and their behavior may deteriorate when encountering transformations insufficiently represented during training (Wan et al.,

2022). Hybrid strategies that use deep learning as an image analyzer while retaining established transform-domain embedding mechanisms can potentially combine the adaptability of neural networks with the structural robustness of classical watermarking (Sy et al., 2020; Zhu et al., 2021).

Optimization techniques offer a complementary solution to scale-factor selection. Metaheuristic algorithms can search continuous parameter spaces more systematically than empirical trial-and-error procedures and have previously been applied to optimize watermark embedding parameters (Sharma et al., 2022). Nevertheless, direct metaheuristic optimization is computationally expensive because each candidate parameter may require complete watermark embedding, attack simulation, extraction, and performance evaluation. Surrogate modeling provides a means of reducing this burden by replacing repeated expensive objective-function evaluations with predictions from a computationally inexpensive approximation model (Bartz-Beielstein, 2021). Accordingly, the present study developed a two-phase hybrid framework integrating U-Net-based adaptive image analysis, regression tree-based surrogate modeling, particle swarm optimization (PSO), and DWT-based watermark embedding. The objective was to improve robustness against geometric attacks while maintaining high perceptual quality and substantially reducing the computational burden associated with scale-factor optimization.

Methods and Materials

The proposed framework consisted of two interconnected phases. In the first phase, a deep CNN based on the U-Net architecture was developed to generate adaptive weight maps identifying image regions suitable for watermark embedding. Six standard 512×512 color images—Baboon, Barbara, Boat, Jet, Lena, and Pepper—were used as principal host images, while a logo image served as the watermark. An augmented dataset containing 2,000 image patches of 128×128 pixels was generated through random cropping, minor rotation, and horizontal and vertical reflection. The dataset was randomly divided into 1,700 training samples (85%) and 300 validation samples (15%).

Ground-truth weight maps were constructed by combining three normalized visual-complexity descriptors: Canny edge information with thresholds of 0.05 and 0.15, local texture information calculated using a 7×7 local-variation window, and Sobel gradient magnitude. Their respective weights were 0.40, 0.30, and 0.30. The resulting maps were Gaussian-smoothed and normalized to $[0, 1]$.

The U-Net contained 35 layers. The encoder used 32- and 64-filter convolutional blocks followed by max pooling, while the bottleneck contained 128-filter convolutional layers. Transposed convolutions and skip connections were used in the decoder to reconstruct spatial information, and a 1×1 convolution produced the final single-channel adaptive weight map. Training employed the Adam optimizer, an initial learning rate of 0.001, mean squared error (MSE) loss, L2 regularization of 0.0001, gradient clipping of 1.0, and mini-batches of 16.

In the second phase, the blue channel of each RGB host image was selected for watermark embedding. A two-level two-dimensional Haar DWT was applied, and the second-level approximation subband (LL2) served as the embedding domain. A regression-tree surrogate model was trained using 30 actual evaluations sampled across the scaling-factor search interval of 0.08–0.26 for each image. PSO subsequently searched the surrogate response surface for 50 iterations to determine an image-specific optimal scaling factor.

The optimized global scaling factor was multiplied pointwise by the CNN-generated local weight map to obtain spatially adaptive embedding strengths. The binary watermark, previously scrambled using

the Arnold transform, was embedded into LL2 coefficients. Inverse DWT reconstructed the modified blue channel, which was recombined with the unchanged red and green channels. Imperceptibility was assessed using PSNR, MSE, and structural similarity index measure (SSIM), whereas robustness was evaluated using normalized correlation (NC) and bit error rate (BER). A combined rotation of 15° and JPEG compression was used as the principal robustness test.

Findings

The U-Net demonstrated highly accurate learning of the adaptive weight maps. Mean validation accuracy reached 99.986%, with a standard deviation of 0.016%, a minimum of 99.944%, and a maximum of 99.994%. Mean MSE was approximately 1.44×10^{-4} , mean RMSE was 0.0108, and mean absolute error was 0.0078. Reconstruction quality was also high, with a mean PSNR of 40.08 dB (SD = 3.38) and mean SSIM of 0.990 (SD = 0.007). Qualitative examination showed that larger weights were concentrated around edges, textures, hair, facial details, object boundaries, and other visually complex regions, whereas smooth areas received substantially lower weights.

The surrogate-assisted PSO successfully identified different optimal scaling factors for different host images, ranging from 0.120 for Lena to 0.147 for Baboon. This variability demonstrated that a single fixed embedding strength was not optimal across images with different structural characteristics. The optimization trajectories stabilized within the 50-iteration search, while the regression-tree models reproduced the nonlinear and image-specific relationships between scaling factor and estimated cost.

Final watermarking performance demonstrated consistently high imperceptibility. PSNR values exceeded 38 dB for all six host images. Lena produced the highest value at approximately 39.45 dB and an SSIM of 0.987, whereas Pepper produced the lowest PSNR at approximately 38.22–38.31 dB while still maintaining acceptable visual quality. Barbara achieved approximately 39.3 dB with an SSIM of 0.966, and Baboon achieved approximately 39.4 dB with an SSIM of 0.976.

Robustness remained high following the combined 15° rotation and JPEG compression attack. NC exceeded 98% for every host image, ranging from approximately 98.09% for Pepper to 98.79% for Barbara. BER remained below 1%, with the lowest value of 0.48% for Barbara and the highest value of 0.76% for Pepper. Baboon achieved an NC of approximately 98.4% and a BER of 0.65%. The extracted watermark remained visually recognizable in all examined images.

The overall mean PSNR of the proposed method was 38.94 dB, while its mean NC reached 98.39%. In comparative evaluation, the NC of the proposed framework exceeded those reported for the examined alternative methods, including approximately 94% for the Zermi method, 92% for DWT-SVD, and 87% for DCT. Although some competing approaches achieved higher PSNR values, the proposed framework demonstrated a stronger preservation of recoverable watermark information under the investigated combined attack.

Discussion and Conclusion

The findings demonstrate that combining adaptive deep visual analysis with surrogate-assisted optimization can effectively address two major weaknesses of conventional transform-domain watermarking: uniform allocation of embedding strength and computationally inefficient parameter selection. The U-Net did not directly perform watermark embedding or extraction; instead, it functioned as an intelligent scene analyzer. This separation allowed the system to exploit deep learning for spatial adaptation without making the entire watermarking process dependent on an end-to-end neural architecture.

The high validation accuracy and near-unity SSIM obtained during weight-map prediction indicate that the CNN successfully learned the intended visual-complexity representation. Concentrating embedding energy in textured regions and edges reduced perceptually detectable distortion in smooth regions while permitting stronger local embedding where modifications could be effectively masked. The image-specific optimal scaling factors further demonstrate that watermark strength should not be treated as a universal constant. Structural differences among images alter their tolerance for embedded information and consequently require adaptive parameterization.

The surrogate model also provided an important computational contribution. Instead of repeatedly executing the complete watermarking, attack, and extraction pipeline during PSO, optimization was conducted primarily on an inexpensive regression-tree approximation of the objective function. This architecture makes extensive exploration of the scaling-factor search space considerably more practical while retaining the ability to model nonlinear cost behavior.

Most importantly, the framework maintained NC values above 98% and BER below 1% under the simultaneous application of rotation and JPEG compression while preserving PSNR above 38 dB. These findings indicate that the method established an effective balance between perceptual transparency and attack robustness. Its principal contribution therefore lies not in any single component, but in the coordinated integration of CNN-based spatial intelligence, surrogate-assisted PSO parameter optimization, and robust DWT-domain embedding.

In conclusion, the proposed framework provides an effective hybrid solution for robust color-image watermarking under geometric and compression attacks. The results support adaptive rather than uniform embedding and demonstrate the practical value of surrogate modeling for computationally expensive watermark optimization. Future development should extend robustness testing to scaling, translation, cropping, row or column removal, and additional signal-processing attacks. Further improvements could include Gaussian-process surrogate models, explicit multi-objective optimization, attack-aware CNN training, and development of a fully blind extraction mechanism that eliminates dependence on the original host image.

نهان نگاری مقاوم تصویر در برابر حملات هندسی با استفاده از بهینه سازی فاکتور مقیاس مبتنی بر مدل جانشین و جاسازی تطبیقی مبتنی بر یادگیری عمیق

مسعود پورمرادی^۱، حمیدرضا غفاری^{۱*}

۱. گروه مهندسی کامپیوتر، واحد فردوس، دانشگاه آزاد اسلامی، فردوس، ایران

*ایمیل نویسنده مسئول: hamidghaffary53@iau.ac.ir

اطلاعات مقاله

چکیده

نوع مقاله

پژوهشی/اصیل

نحوه استناد به این مقاله:

پورمرادی، مسعود، و غفاری، حمیدرضا. (۱۴۰۵). نهان نگاری مقاوم تصویر در برابر حملات هندسی با استفاده از بهینه سازی فاکتور مقیاس مبتنی بر مدل جانشین و جاسازی تطبیقی مبتنی بر یادگیری عمیق. علم تصمیم گیری و سیستم های هوشمند، ۳(۲)، ۱-۳۷.



© ۱۴۰۵ تمامی حقوق انتشار این مقاله متعلق به نویسنده است. انتشار این مقاله به صورت دسترسی آزاد مطابق با گواهی (CC BY-NC 4.0) صورت گرفته است.

نهان نگاری دیجیتال به عنوان راهکاری بنیادین برای حفاظت از حق نشر، همواره با چالش ذاتی برقراری تعادل میان شفافیت و مقاومت، به ویژه در برابر حملات هندسی، روبرو است. این تعادل حساس به شدت تحت تأثیر دو عامل کلیدی قرار دارد: انتخاب بهینه فاکتور مقیاس (قدرت جاسازی) و انتخاب هوشمندانه ی نواحی جاسازی. این مقاله یک چارچوب نهان نگاری ترکیبی و مقاوم برای تصاویر رنگی ارائه می دهد که با تلفیق راهبردی تحلیل هوشمند تصویر و بهینه سازی تسریع شده، به طور خاص برای مقابله با حملات چرخش طراحی شده است. روش پیشنهادی دو نوآوری اصلی را در خود ادغام می کند. نخست، یک شبکه ی عصبی کانولوشنی عمیق با معماری U-Net برای تولید یک نقشه ی وزن تطبیقی آموزش داده می شود. این شبکه، مستقل از فرآیند جاسازی، محتوای تصویر میزبان را تحلیل کرده و به طور خودکار نواحی با پیچیدگی بصری بالا (لبه ها و بافت ها) را شناسایی می کند. دوم، یک الگوریتم بهینه سازی ازدحام ذرات (PSO) که توسط یک مدل جانشین مبتنی بر درخت رگرسیون هدایت می شود، برای یافتن مقدار بهینه فاکتور مقیاس به کار می رود. این رویکرد با جایگزینی ارزیابی های پرهزینه ی تابع واقعی با یک تخمین گر سریع، شتاب چشمگیری به فرآیند بهینه سازی می بخشد. در نهایت، یک جاسازی کاملاً تطبیقی با تلفیق نقطه ای فاکتور مقیاس بهینه و نقشه ی وزن محلی در زیرباند تقریبی تبدیل موجک گسسته (DWT) انجام می شود. نتایج تجربی بر روی تصاویر استاندارد، کارایی برجسته ی روش را تأیید می کنند. این روش با دستیابی به میانگین PSNR بالای ۳۸ دسی بل و ضریب همبستگی نرمال شده (NC) بالای ۰.۹۸ پس از اعمال هم زمان حمله ی چرخش و فشرده سازی JPEG، پیشرفت قابل توجهی را نسبت به روش های پیشین به نمایش گذاشته و تعادل میان شفافیت و مقاومت را به نحوی مؤثر برقرار می سازد.

کلیدواژگان: نهان نگاری تصویر، حملات هندسی، بهینه سازی فاکتور مقیاس، مدل جانشین، یادگیری عمیق، شبکه عصبی کانولوشنی

۱. مقدمه

در عصر ارتباطات دیجیتال، تولید، توزیع و مصرف محتوای چندرسانه‌ای به طرز بی‌سابقه‌ای افزایش یافته است. تصاویر دیجیتال به عنوان یکی از مهم‌ترین ارکان این محتوا، در بسترهای گوناگونی چون شبکه‌های اجتماعی، وبسایت‌های خبری، پایگاه‌های داده پزشکی و سیستم‌های تجارت الکترونیک به اشتراک گذاشته می‌شوند [۱]. با این حال، همین سهولت در تبادل، چالش‌های امنیتی جدی را به همراه دارد. امکان کپی‌برداری غیرمجاز، دستکاری محتوا^۱ و توزیع بدون ذکر منبع، ناقضان حق نشر^۲ و اعتبارسنجی محتوا را به تهدیداتی همیشگی تبدیل کرده است [۲]. در این میان، روش‌های سنتی رمزنگاری^۳ اگرچه می‌توانند از محتوای دیجیتال در حین انتقال محافظت کنند، اما پس از رمزگشایی، هیچ سازوکاری برای جلوگیری از کپی و توزیع غیرمجاز یا اثبات مالکیت ارائه نمی‌دهند. نهان‌نگاری دیجیتال^۴ به عنوان یک فناوری مکمل و قدرتمند، با جاسازی یک نشانه دیجیتال^۵ به صورت نامحسوس درون رسانه میزبان^۶، راه‌حلی پایدار برای احراز هویت، اثبات مالکیت و ردیابی توزیع غیرقانونی فراهم می‌آورد [۳].

یک سیستم نهان‌نگاری ایده‌آل، حاصل یک مصالحه^۷ پیچیده میان سه ویژگی بنیادین و اغلب متعارض است: شفافیت^۸، مقاومت^۹ و ظرفیت^{۱۰} [۴]. شفافیت ایجاب می‌کند که جاسازی واترمارک نباید هیچ‌گونه افت کیفیت بصری محسوسی در تصویر میزبان ایجاد کند. مقاومت به معنای توانایی الگوریتم در بازیابی موفق واترمارک حتی پس از آنکه تصویر میزبان مورد حملات عمدی یا غیرعمدی پردازش سیگنال قرار گرفته باشد، است. ظرفیت نیز به میزان اطلاعاتی که می‌توان در یک رسانه جاسازی کرد، اشاره دارد. در این میان، مقاومت در برابر حملات هندسی^{۱۱}، به ویژه حملات معروف به RST (چرخش، تغییر مقیاس و انتقال^{۱۲})، یکی از بزرگ‌ترین چالش‌ها در طراحی سیستم‌های نهان‌نگاری است [۵]. این حملات با ایجاد ناهم‌زمانی^{۱۳} بین مراحل جاسازی و استخراج، بدون آنکه نیازی به تخریب مستقیم واترمارک داشته باشند، الگوریتم استخراج را در مکانیابی و بازیابی نشانه دیجیتال ناتوان می‌سازند. فرآیند نهان‌نگاری تصویر به طور کلی شامل دو مرحله اصلی جاسازی^{۱۴} و استخراج^{۱۵} است. در مرحله جاسازی، واترمارک در تصویر میزبان پنهان می‌شود و در مرحله استخراج، واترمارک از تصویر نهان‌نگاری شده بازیابی می‌گردد. یک نمایش کلی از این فرآیند، شامل رمزنگاری واترمارک پیش از جاسازی برای افزایش امنیت، در پژوهش [۶] ارائه شده است.

راهکارهای نهان‌نگاری مقاوم را می‌توان به طور کلی به دو دسته حوزه مکان^{۱۶} و حوزه تبدیل^{۱۷} تقسیم‌بندی کرد. روش‌های حوزه مکان که با دستکاری مستقیم مقادیر پیکسل‌ها کار می‌کنند، ساده و سریع هستند اما معمولاً مقاومت کمی در برابر حملات دارند. در مقابل،

¹ Tampering² Copyright³ Cryptography⁴ Digital Watermarking⁵ Watermark⁶ Host⁷ Trade-off⁸ Imperceptibility⁹ Robustness¹⁰ Capacity¹¹ Geometric Attacks¹² Rotation, Scaling, and Translation¹³ Desynchronization¹⁴ Embedding¹⁵ Extracting¹⁶ Spatial Domain¹⁷ Transform Domain

روش‌های حوزه تبدیل، ابتدا تصویر را توسط یک تبدیل ریاضی مانند تبدیل کسینوسی گسسته^(DCT¹)، تبدیل موجک گسسته^(DWT²) یا تجزیه مقادیر منفرد^(SVD³) به فضای فرکانسی می‌برند و سپس واترمارک را در ضرایب این فضاها جاسازی می‌کنند. این رویکرد به دلیل تطابق بهتر با سیستم بینایی انسان^(HVS⁴) و توزیع هوشمندانه‌تر انرژی واترمارک، به طور بالقوه مقاومت و شفافیت بالاتری را ارائه می‌دهد [7].

با این حال، موفقیت نهایی یک الگوریتم حوزه تبدیل، به شدت به دو تصمیم کلیدی وابسته است. نخست، انتخاب دقیق و بهینه فاکتور مقیاس (α) است که قدرت و انرژی واترمارک جاسازی شده را تعیین می‌کند. یک مقدار α بسیار کوچک، اگرچه شفافیت را در سطح عالی نگه می‌دارد، اما مقاومت لازم برای بقا در برابر حملات را فراهم نمی‌کند. برعکس، یک α بسیار بزرگ، مقاومت را افزایش می‌دهد اما با ایجاد اعوجاج قابل مشاهده، شفافیت تصویر را از بین می‌برد [8]. دومین چالش اساسی، مکان‌یابی هوشمندانه برای جاسازی است. روش‌های سنتی اغلب واترمارک را به صورت یکنواخت در تمام نواحی یک زیرباند فرکانسی (مثلاً LL) توزیع می‌کنند. این در حالی است که نواحی مختلف تصویر، ظرفیت پنهان‌سازی متفاوتی دارند. نواحی صاف و یکنواخت، کوچک‌ترین تغییر را به سرعت برای بیننده آشکار می‌کنند، در حالی که نواحی پر جزئیات، لبه‌ها و بافت‌های پیچیده، می‌توانند به طور طبیعی‌تری نویز جاسازی را استتار کنند. بنابراین، یک سیستم کارا باید قادر به تشخیص این نواحی و تخصیص قدرت جاسازی به صورت تطبیقی باشد [9].

در سال‌های اخیر، ظهور و بلوغ فناوری‌های یادگیری عمیق⁵، به‌ویژه شبکه‌های عصبی کانولوشنی^(CNNs⁶)، افق‌های جدیدی را در تحلیل و درک محتوای تصویر گشوده است [9]. این شبکه‌ها با توانایی خارق‌العاده خود در استخراج سلسله‌مراتبی ویژگی‌های بصری از لبه‌های سطح پایین تا مفاهیم پیچیده، ابزاری ایده‌آل برای پرداختن به چالش دوم، یعنی مکان‌یابی هوشمندانه جاسازی، فراهم می‌کنند. از سوی دیگر، چالش انتخاب بهینه فاکتور مقیاس، یک مسئله بهینه‌سازی پیوسته است که حل آن با روش‌های سنتی سعی و خطا غیرممکن است. الگوریتم‌های فراابتکاری⁷ مانند بهینه‌سازی ازدحام ذرات^(PSO⁸) راه‌حلی مؤثر برای کاوش فضای جستجو ارائه می‌دهند [10]. با این وجود، استفاده مستقیم از این الگوریتم‌ها در مسائل نهان‌نگاری بسیار پرهزینه است، چرا که تابع هدف (مقاومت/شفافیت) برای هر ارزیابی، نیازمند اجرای کامل فرآیند جاسازی، اعمال حمله و استخراج واترمارک بر روی یک مجموعه داده بزرگ است. برای غلبه بر این مانع، مفهوم مدل‌های جانشین⁹ که تقریب‌های سریع و کم‌هزینه از تابع هدف اصلی ارائه می‌دهند، مطرح شده است [11].

در این پژوهش، ما یک چارچوب نهان‌نگاری ترکیبی و هوشمند ارائه می‌دهیم که با تلفیق استراتژیک پیشرفت‌های حوزه‌های مختلف، به طور هم‌زمان به هر دو چالش فوق پاسخ می‌دهد. نوآوری‌های کلیدی این مقاله عبارتند از:

۱. ارائه یک سیستم دو مرحله‌ای تلفیقی: معماری پیشنهادی شامل یک فاز تحلیل مبتنی بر CNN و یک فاز جاسازی بهینه‌سازی شده است که نقاط قوت یادگیری عمیق و روش‌های حوزه تبدیل را با یکدیگر ترکیب می‌کند.

¹ Discrete Cosine Transform

² Discrete Wavelet Transform

³ Singular Value Decomposition

⁴ Human Visual System

⁵ Deep Learning

⁶ Convolutional Neural Networks

⁷ Metaheuristic

⁸ Particle Swarm Optimization

⁹ Surrogate Models

۲. طراحی یک شبکه CNN پیشگو و مستقل از جاسازی: ما یک شبکه U-Net¹ سبک و کارا را برای تحلیل ساختاری تصویر میزبان و تولید یک نقشه وزن² تطبیقی آموزش می دهیم. این نقشه، شایستگی هر ناحیه از تصویر برای جاسازی مقاوم و شفاف را نشان می دهد. بر خلاف بسیاری از روش های نهان نگاری عمیق سرتاسری^۳، شبکه ما مستقیماً در فرآیند جاسازی/استخراج دخالت ندارد و صرفاً به عنوان یک تحلیل گر صحنه عمل می کند.

۳. معرفی یک فرآیند بهینه سازی مبتنی بر مدل جانشین: برای اولین بار، با موفقیت از یک مدل جانشین مبتنی بر درخت رگرسیون برای تخمین تابع هزینه در بهینه سازی فاکتور مقیاس استفاده می شود. این مدل که بر روی تعداد محدودی نمونه آموزش دیده، جایگزین ارزیابی های پرهزینه تابع واقعی شده و بهینه سازی PSO را به طور چشمگیری تسریع می کند.

۴. جاسازی کاملاً تطبیقی: فاکتور مقیاس بهینه یافته سراسری (α) با نقشه وزن محلی CNN تلفیق می شود تا یک قدرت جاسازی نقطه های و تطبیقی ایجاد کند. این امر تضمین می کند که انرژی واترمارک به صورت هوشمندانه در نواحی با پتانسیل پنهان سازی بالا متمرکز شده و از آسیب رسانی به نواحی حساس تصویر جلوگیری می شود.

ساختار ادامه مقاله به شرح زیر است: در بخش ۲، پیشینه جامعی از پژوهش های مرتبط با روش پیشنهادی ارائه می شود. بخش ۳ به تشریح دقیق و گام به گام روش پیشنهادی، از تولید نقشه وزن CNN تا جاسازی بهینه شده، اختصاص دارد. نتایج تجربی، ارزیابی های کمی و کیفی و مقایسه با روش های پیشین در بخش ۴ ارائه می شوند. در نهایت، بخش ۵ به نتیجه گیری و ارائه پیشنهادات برای پژوهش های آتی می پردازد.

۲- مروری جامع بر پیشینه پژوهش

حوزه نهان نگاری تصویر، به ویژه با هدف دستیابی به مقاومت در برابر حملات هندسی، یکی از پویاترین زمینه های تحقیقاتی در پردازش سیگنال و امنیت چندرسانه ای است. تلاش های پژوهشی را می توان در قالب یک سیر تکاملی از روش های کلاسیک حوزه تبدیل به سمت رویکردهای هوشمند مبتنی بر بهینه سازی و در نهایت، یادگیری عمیق دسته بندی کرد. در این بخش، مروری عمیق بر این پیشینه ارائه می شود تا جایگاه و نوآوری پژوهش حاضر به روشنی تبیین گردد.

۲-۱- روش های پایه در حوزه تبدیل: از DCT تا SVD

نسل اول سیستم های نهان نگاری مقاوم، عمدتاً بر روی استفاده از تبدیلات حوزه فرکانس برای ارائه یک مصالحه اولیه بین شفافیت و مقاومت متمرکز بودند.

- روش های مبتنی بر DCT: تبدیل کسینوسی گسسته به دلیل خاصیت تراکم انرژی⁴ عالی و تطابق نسبی با استاندارد فشرده سازی JPEG، یکی از محبوب ترین انتخاب ها بوده است. یک نقطه عطف در این حوزه، کار Cox و همکاران (۱۹۹۷) [۱۲] بود که اساس طیف گسترده⁵ را برای نهان نگاری بنا نهادند و واترمارک را در ضرایب میانی DCT جاسازی کردند. در ادامه، [۱۳] یک تکنیک DCT مقاوم در برابر JPEG ارائه دادند که با مفهوم باقیمانده ریاضی به تنظیم دقیق ضرایب فرکانس پایین می پرداخت [۱۴]. با ارائه روش DCT دو سطحی بر روی تصاویر رنگی، گامی فراتر نهاده و واترمارک رنگی را در ضرایب DC و AC جاسازی

¹ U-shaped Network

² Weight Map

³ deep end-to-end watermarking

⁴ Energy Compaction

⁵ Spread Spectrum

کردند که به مقاومت بالایی در برابر حملات رایج دست یافت. با این حال، ماهیت سراسری DCT و عدم تفکیک فرکانسی-مکانی، انعطاف‌پذیری این روش‌ها را برای جاسازی تطبیقی و مقاومت در برابر حملات هندسی محدود می‌کرد.

- روش‌های مبتنی بر DWT: تبدیل موجک گسسته با ارائه یک نمایش چندوضوحی¹ از تصویر، توانست بر محدودیت DCT در تحلیل هم‌زمان فرکانس و مکان غلبه کند. این تبدیل با تجزیه تصویر به زیرباندهای تقریبی (LL) و جزئیات (LH, HL, HH)، امکان انتخاب زیرباند مناسب برای جاسازی را فراهم می‌کند [۱۵]. یک روش جامع با تجزیه موجک سه سطحی در فضای رنگی غیرهمبسته ارائه دادند و با کمک الگوریتم کلونی زنبور مصنوعی (ABC)، فاکتور مقیاس بهینه برای جاسازی در تمامی زیرباندها را یافتند [۱۶]. نیز یک روش بلوک‌بندی مبتنی بر DWT برای تصاویر رنگی پیشنهاد دادند که بر مقاومت در برابر نویز و فشرده‌سازی تمرکز داشت.

- روش‌های مبتنی بر SVD و تجزیه‌های ماتریسی: SVD به دلیل پایداری ذاتی مقادیر تکین² در برابر تغییرات جزئی و برخی حملات هندسی، به ابزاری قدرتمند تبدیل شد. روش‌های مبتنی بر SVD، مانند کار [17] برای تصاویر پزشکی، واترمارک را با اصلاح مقادیر تکین تصویر یا زیرباندهای آن جاسازی می‌کنند. با این حال، کاربرد مستقل SVD ممکن است با مشکل مثبت کاذب³ در استخراج مواجه شود. با این حال، کاربرد مستقل SVD ممکن است با مشکل مثبت کاذب در استخراج مواجه شود. برای رفع این مشکل و نیز کاهش هزینه محاسباتی بالای SVD، نسخه‌های توسعه‌یافته‌ای مانند RSVD (SVD افزونه) پیشنهاد شده است [18]. در بسیاری از طرح‌های واترمارکینگ، بکارگیری تکنیک‌هایی همچون ممان زرنایک⁴ نیز می‌تواند توانایی مقاومت در برابر حملات هندسی و نویز را تقویت کند. راه حل دیگر رفع این مشکل و بهره‌گیری از مزایای هم‌افزایی، روش‌های ترکیبی مانند DWT-DCT-SVD هستند [۷]. یک نمونه شاخص از این رویکرد است که واترمارک را پس از اعمال متوالی DCT و DWT بر روی تصویر و تقسیم واترمارک، جاسازی می‌کند. نقطه ضعف مشترک تمامی این روش‌های ارزشمند، نگاه ایستا به تصویر و انتخاب تجربی فاکتور مقیاس است. آن‌ها تصویر را به صورت یک موجودیت یکپارچه می‌بینند و از ظرفیت پنهان‌سازی متغیر در نواحی مختلف آن چشم‌پوشی می‌کنند.

۲-۲- روش‌های مبتنی بر بهینه‌سازی و ضرورت مدل‌های جانشین

برای غلبه بر محدودیت انتخاب تجربی فاکتور مقیاس، الگوریتم‌های بهینه‌سازی فراابتکاری به عنوان راهکاری مؤثر وارد میدان شدند. این الگوریتم‌ها، مسئله نهان‌نگاری را به یک مسئله بهینه‌سازی تبدیل می‌کنند که در آن، هدف، یافتن پارامترهایی (عمدتاً فاکتور مقیاس و گاهی مکان‌های جاسازی) است که یک تابع هزینه چندهدفه (ترکیب وزنی از معیارهای شفافیت و مقاومت) را کمینه کنند. محققان در [۱۰] یک رویکرد مبتنی بر DWT-SVD-HD ارائه دادند که در آن، فاکتور مقیاس توسط الگوریتم کرم شبتاب (FA⁵) بهینه‌سازی می‌شد. این روش با در نظر گرفتن هم‌زمان استحکام و نامحسوس بودن، بهبود قابل توجهی را نسبت به روش‌های با پارامترهای دستی نشان داد. [19] نیز از تبدیل آرنولد برای رمزنگاری واترمارک و SVD برای جاسازی بهینه‌شده استفاده کردند. [۲۰] یک روش واترمارک چندگانه در حوزه DCT ارائه دادند که در آن، کدهای تکرار شونده برای افزایش مقاومت به کار رفته بود.

¹ Multi-resolution

² Singular Values

³ False Positive

⁴ Zernike Moments

⁵ Firefly Algorithm

با این حال، یک چالش اساسی در کاربرد مستقیم الگوریتم های فراابتکاری وجود دارد: هزینه محاسباتی گزاف. در هر تکرار از الگوریتم بهینه ساز (برای هر ذره در PSO یا هر کرم شبتاب در FA)، یک کاندیدای جدید برای فاکتور مقیاس باید ارزیابی شود. این ارزیابی شامل جاسازی واترمارک با آن فاکتور، اعمال یک یا چند حمله به تصویر، استخراج واترمارک و محاسبه نهایی خطا (مثلاً BER) است. این فرآیند برای یک مجموعه داده تست، فوق العاده زمان بر بوده و عملاً استفاده از الگوریتم های با جمعیت بالا و تکرار زیاد را غیرممکن می سازد. راه حلی که در دیگر شاخه های مهندسی و اخیراً در یادگیری ماشین برای تنظیم هایپر پارامترها مرسوم شده، استفاده از مدل های جانشین است [۱۱]. یک مدل جانشین، یک مدل تقریبی و سریع (مانند رگرسیون، درخت تصمیم یا فرآیند گوسی) است که بر روی تعداد محدودی ارزیابی واقعی از تابع هدف آموزش می بیند و سپس می تواند به عنوان یک تخمین گر کم هزینه برای هدایت الگوریتم بهینه ساز به کار رود. این دقیقاً همان شکاف پژوهشی است که ما در این مقاله به آن می پردازیم.

۲-۳- ظهور یادگیری عمیق در نهان نگاری

در سال های اخیر، یادگیری عمیق، انقلابی در پردازش تصویر ایجاد کرده و به طور طبیعی، کاربردهای آن در نهان نگاری نیز مورد توجه شدید قرار گرفته است. مرور جامع [۲] نشان می دهد که روش های مبتنی بر یادگیری عمیق به سرعت در حال تبدیل شدن به جریان اصلی هستند.

- روش های انتها به انتها^۱: این دسته از روش ها، کل فرآیند جاسازی و استخراج را به عنوان یک شبکه عصبی عمیق مدل سازی می کنند. یک شبکه رمزگذار (مسیر انقباضی)، تصویر میزبان و واترمارک را دریافت کرده و تصویر جاسازی شده را تولید می کند. یک شبکه رمزگشا (مسیر انبساطی) نیز وظیفه استخراج واترمارک را بر عهده دارد. اگرچه این روش ها نتایج چشمگیری در شفافیت ارائه می دهند، اما معمولاً مقاومت کمتری در برابر حملات هندسی دارند، مگر آنکه در حین آموزش، نمونه های حملات دیده شده^۲ به شبکه خورنده شوند که این خود می تواند فرآیند آموزش را ناپایدار و بسیار طولانی کند [2].
- روش های تلفیقی^۳: این رویکرد که پژوهش ما نیز در این دسته قرار می گیرد، از قدرت تحلیل CNN ها در کنار ساختارهای کلاسیک و مقاوم استفاده می کند. [9] یک چارچوب قدرتمند به نام MDResNet-HDWM ارائه دادند. در این روش، یک شبکه ResNet چند مقیاسی، نقاط کلیدی مقاوم در برابر حملات هندسی را در تصویر شناسایی کرده و واترمارک تنها در نواحی اطراف این نقاط جاسازی می شود. این روش به دلیل گره زدن فرآیند استخراج به نقاط کلیدی، مقاومت ذاتی بسیار بالایی در برابر حملات RST دارد. [21] نیز یک روش نهان نگاری کور مبتنی بر CNN در حوزه موجک ارائه دادند که در آن، از یک شبکه کانولوشنی برای یادگیری رابطه بین ضرایب موجک و بیت های واترمارک استفاده شد.

ایده کلیدی که از پژوهش های فوق الهام گرفته می شود، این است که CNN ها تحلیل گرهای صحنه بی نظیری هستند. به جای آنکه از یک شبکه برای کل فرآیند استفاده کنیم، می توانیم از یک شبکه سبک تر و هدفمندتر صرفاً برای تولید یک نقشه راهنمای جاسازی استفاده کنیم. این نقشه می تواند پیچیدگی های بصری را کدگذاری کرده و یک فرآیند جاسازی کلاسیک اما اثبات شده در حوزه موجک را به صورت تطبیقی هدایت کند. این رویکرد، مزایای هر دو جهان را دارد: مقاومت ساختاری روش های حوزه موجک و هوشمندی و تطبیق پذیری CNN ها.

¹ End-to-End

² Augmented Data

³ Hybrid Deep Learning

۲-۴- جمع بندی و جایگاه پژوهش حاضر

با بررسی جامع پیشینه، شکاف های پژوهشی زیر به وضوح قابل مشاهده است:

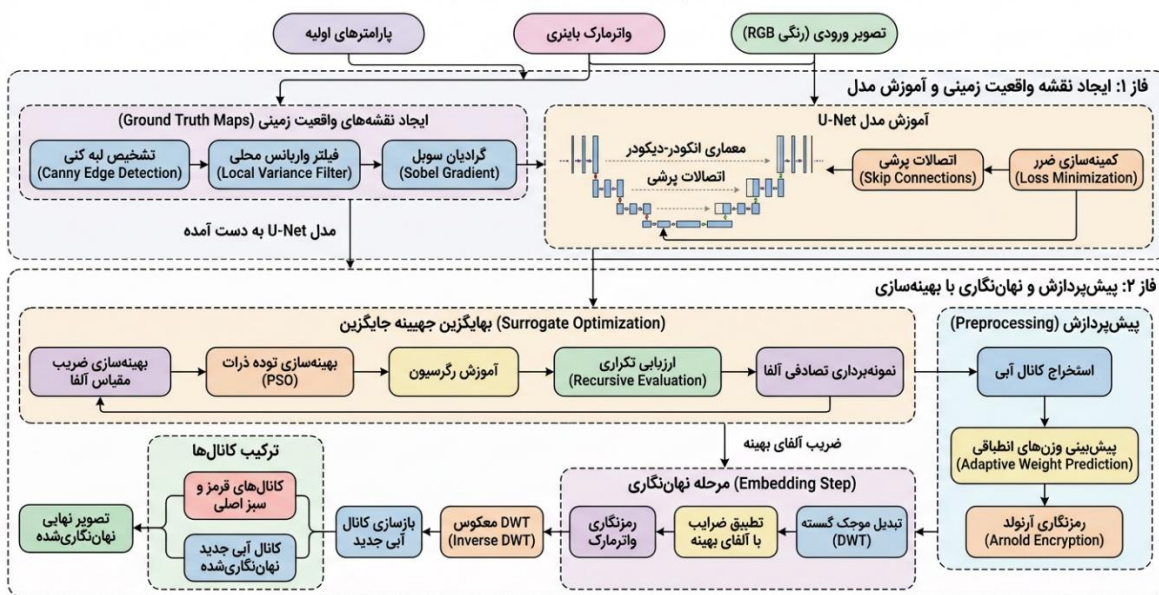
۱. عدم تطبیق پذیری کامل: روش های کلاسیک حوزه تبدیل، فاقد مکانیزمی هوشمند برای تشخیص و بهره برداری از ظرفیت پنهان سازی متغیر در نواحی مختلف تصویر هستند.
۲. هزینه بالای بهینه سازی مستقیم: کاربرد مستقیم الگوریتم های فراابتکاری برای بهینه سازی پارامترهای نهان نگاری، به دلیل ارزیابی های مکرر و پرهزینه تابع واقعی، عملاً ناکارآمد است.
۳. پیچیدگی آموزش نهان نگاری عمیق سرتاسری: روش های تماماً مبتنی بر یادگیری عمیق، برای دستیابی به مقاومت در برابر حملات هندسی، نیازمند آموزش های پیچیده و طولانی با مجموعه داده های عظیم هستند و اغلب مقاومتشان در برابر حملات دیده نشده تضمین شده نیست.

پژوهش حاضر دقیقاً در محل تلاقی این شکاف ها قرار می گیرد و یک راه حل یکپارچه و نوآورانه ارائه می دهد. ما برای اولین بار، یک مدل جانشین را برای تسریع بهینه سازی فاکتور مقیاس در کنار یک شبکه CNN که برای تولید نقشه وزن تطبیقی آموزش دیده، ترکیب می کنیم. این تلفیق منحصربه فرد، یک سیستم نهان نگاری را شکل می دهد که نه تنها از نظر مفهومی نوین است، بلکه از نظر عملی نیز کارایی بالایی در برقراری تعادل میان شفافیت و مقاومت در برابر حملات هندسی دارد.

۳- روش پیشنهادی: چارچوب نهان نگاری ترکیبی مبتنی بر CNN و مدل جانشین

در این بخش، جزئیات فنی و مبانی نظری چارچوب پیشنهادی ارائه می گردد. روش ارائه شده یک سیستم دو فازه است که هر فاز آن برای رفع یکی از چالش های بنیادین در نهان نگاری مقاوم طراحی شده است: فاز اول به تحلیل هوشمند محتوای تصویر برای شناسایی نواحی مستعد جاسازی می پردازد و فاز دوم، با بهره گیری از یک مکانیزم بهینه سازی سریع، فرآیند جاسازی تطبیقی را به انجام می رساند. شمای مفهومی و گردش کار کلی این چارچوب در شکل ۱ به نمایش درآمده است.

نمودار جریان پردازش تصویر و نهان نگاری انطباقی



شکل افلوچارت روش پیشنهادی

همان طور که در شکل ۱ مشاهده می شود، خط لوله پردازشی روش پیشنهادی از دو فاز مجزا اما کاملاً مرتبط تشکیل شده است. نقطه شروع، تصویر ورودی رنگی (RGB) و واترمارک باینری است که به موازات یکدیگر وارد سیستم می شوند. در فاز اول، هدف اصلی، ساخت یک تحلیل گر بصری هوشمند است. این تحلیل گر که مبتنی بر یک شبکه عصبی کانولوشنی با معماری U-Net پیاده سازی شده، ابتدا با استفاده از یک مجموعه داده آموزشی که شامل جفت های تصویر میزبان و نقشه حقیقت مبنا است، آموزش می بیند. نقشه های حقیقت مبنا^۱ از پیش و با ترکیب هوشمندانه ای از معیارهای بینایی ماشین نظیر آشکارساز لبه کنی^۲، تحلیل بافت و محاسبه بزرگی گرادیان تولید می شوند. فرآیند آموزش، شامل کمینه سازی خطای بازسازی^۳ در معماری Encoder-Decoder مجهز به اتصالات پرشی^۴ است. خروجی نهایی این فاز، یک مدل آموزش دیده است که قادر است برای هر تصویر دلخواه، یک نقشه وزن تطبیقی تولید کند. این نقشه، یک ماتریس دو بُعدی است که در آن، مقادیر بالاتر نشان دهنده نواحی با پیچیدگی بصری بیشتر (مانند لبه ها و بافت ها) و در نتیجه، مستعدتر برای جاسازی مقاوم و نامحسوس واترمارک است.

فاز دوم، هسته عملیاتی سیستم را تشکیل می دهد و خود به دو مسیر موازی تقسیم می شود که در نهایت با یکدیگر تلفیق می گردند. در این فاز، تصویر میزبان اصلی ابتدا وارد مرحله پیش پردازش می شود که مهم ترین گام آن، جداسازی و استخراج کانال آبی تصویر است. این انتخاب مبتنی بر حساسیت کمتر سیستم بینایی انسان به تغییرات در این طیف رنگی است. هم زمان، پارامترهای اولیه سیستم از جمله فاکتور مقیاس، وارد یک حلقه بهینه سازی می شوند. این حلقه با محوریت مدل جانشین و الگوریتم PSO، مقدار بهینه فاکتور مقیاس را به نحوی تعیین می کند که تعادل میان شفافیت و مقاومت را برقرار سازد. استفاده از مدل جانشین در این مرحله، یک نوآوری کلیدی برای کاهش چشمگیر هزینه های محاسباتی ناشی از ارزیابی های مکرر است.

پس از تعیین فاکتور مقیاس بهینه و دریافت نقشه وزن از فاز اول، هر دو به مرحله نهان نگاری هدایت می شوند. در این گام، جاسازی واترمارک باینری به صورت تطبیقی و نقطه ای در ضرایب زیرباند تقریبی (LL) حاصل از تبدیل موجک گسسته انجام می پذیرد. قدرت جاسازی در هر پیکسل، حاصل ضرب فاکتور مقیاس بهینه سراسری در وزن محلی تعیین شده توسط CNN است. این مکانیزم تضمین می کند که انرژی واترمارک به طور هوشمندانه در نواحی پر جزئیات تصویر متمرکز شده و از نواحی صاف و حساس دوری می کند. در نهایت، با بازسازی کانال آبی جدید از طریق اعمال معکوس تبدیل موجک و ادغام آن با کانال های قرمز و سبز دست نخورده، تصویر نهایی نهان نگاری شده حاصل می شود. این تصویر باید به گونه ای باشد که هم کیفیت بصری بالایی داشته باشد و هم در برابر حملات هندسی و پردازشی مقاومت کند. در ادامه این بخش، مبانی نظری و جزئیات ریاضی هر یک از این مراحل به طور کامل تشریح می گردد.

• ۳-۱- فاز اول: تولید نقشه وزن تطبیقی مبتنی بر شبکه U-Net

هدف بنیادین در این فاز، طراحی یک تحلیل گر صحنه هوشمند است که بتواند فارغ از فرآیند جاسازی، محتوای تصویر میزبان را درک کرده و یک راهنمای مکانی برای تخصیص قدرت جاسازی تولید کند. این راهنما که آن را نقشه وزن^۵ می نامیم، یک ماتریس دو بُعدی با ابعاد تصویر میزبان است که در آن، هر پیکسل مقداری بین صفر (کاملاً نامناسب برای جاسازی) و یک (ایده آل برای جاسازی) دارد. برای دستیابی به این هدف، یک شبکه عصبی کانولوشنی عمیق با معماری U-Net طراحی و پیاده سازی شده است. انتخاب این معماری به دلیل

¹ Ground Truth Maps

² Canny

³ Loss Minimization

⁴ Skip Connections

⁵ Weight Map

توانایی ذاتی آن در استخراج سلسله مراتبی ویژگی ها از طریق مسیر انقباضی^۱ و بازیابی دقیق جزئیات مکانی از طریق مسیر انبساطی^۲ و اتصالات پرشی^۳ صورت گرفته است. این ویژگی ها برای تولید یک نقشه وزنی با وضوح کامل که کوچک ترین جزئیات ساختاری تصویر را منعکس کند، حیاتی است.

۳-۱-۱- طراحی و تولید داده های هدف

از آنجا که نقشه های وزن ایده آل به صورت آماده در دسترس نیستند، یک گام اساسی در این فاز، تولید یک مجموعه داده آموزشی با کیفیت بالا شامل جفت های تصویر میزبان و نقشه وزن هدف است. برای ساخت نقشه هدف، که به عنوان حقیقت مبنا برای آموزش شبکه عمل می کند، یک روش ترکیبی مبتنی بر اصول سیستم بینایی انسان (HVS^۴) طراحی شده است. منطق اساسی این است که نواحی با پیچیدگی بصری بالا، مانند لبه ها و بافت های متراکم، به دلیل اثر پوشاندگی^۵، تغییرات ناشی از جاسازی واترمارک را بهتر استتار می کنند. بر این اساس، برای هر تصویر در مجموعه داده آموزشی، نقشه حقیقت مبنا (G) از ترکیب وزنی سه معیار بینایی محاسباتی زیر ساخته می شود:

نقشه لبه^۶ (E): برای استخراج لبه های قوی و شاخص تصویر، از آشکارساز لبه کنی با آستانه های [۰.۰۵, ۰.۱۵] استفاده می شود. خروجی این مرحله یک نقشه باینری است که در آن، پیکسل های لبه مقدار ۱ و سایر پیکسل ها مقدار ۰ دارند. لبه ها به عنوان شاخص ترین نشانه های بصری، بهترین نامزدها برای جاسازی مقاوم هستند.

نقشه بافت^۷ (T): مطابق با رابطه (۱) تحلیل بافت به ما در یافتن نواحی با تغییرات شدت نور بالا و تکرار شونده کمک می کند. برای کمی سازی این مفهوم، از فیلتر واریانس محلی^۸ بر روی تصویر استفاده می شود. یک پنجره با ابعاد مشخص (مثلاً ۷×۷) بر روی تصویر حرکت کرده و واریانس مقادیر پیکسل های درون آن محاسبه می شود. مقادیر بالای واریانس، نشان دهنده بافت های پیچیده و مناسب برای جاسازی است.

$$T(x, y) = \frac{1}{|W|} \sum_{(i, j) \in W_{x, y}} (I(i, j) - \mu W_{x, y})^2 \quad (1)$$

که در آن $I(i, j)$ شدت روشنایی پیکسل، $W_{x, y}$ پنجره اطراف پیکسل (x, y) ، $|W|$ ، تعداد پیکسل های پنجره، و $\mu W_{x, y}$ میانگین شدت روشنایی در آن پنجره است. خروجی نهایی نرمال سازی می شود.

نقشه بزرگی گرادیان^۹ (M): این معیار بر قدرت لبه ها تأکید دارد. بزرگی گرادیان تصویر با استفاده از عملگرهای مشتق افقی G_x و عمودی G_y (مانند فیلتر سوبل^{۱۰}) محاسبه شده و سپس اندازه بردار گرادیان برای هر پیکسل مطابق با رابطه (۲) به دست می آید.

$$M(x, y) = \sqrt{G_x(x, y)^2 + G_y(x, y)^2} \quad (2)$$

در نهایت، نقشه حقیقت مبنا (G) از ترکیب خطی این سه معیار نرمال سازی شده مطابق با رابطه (۳) ایجاد می شود:

¹ Encoder

² Decoder

³ Skip Connections

⁴ Human Visual System

⁵ Masking Effect

⁶ Edge Map

⁷ Texture Map

⁸ Local Variance

⁹ Gradient Magnitude Map

¹⁰ Sobel

$$G = \lambda_1 E_{norm} + \lambda_2 T_{norm} + \lambda_3 M_{norm} \quad (3)$$

که در آن $\lambda_1, \lambda_2, \lambda_3$ اوزان تجربی هستند که اهمیت نسبی هر معیار را مشخص می‌کنند و تابع شرط $\lambda_1 + \lambda_2 + \lambda_3 = 1$ است. در پیاده‌سازی ما، وزن‌های 0.4, 0.3, 0.3 به ترتیب برای لبه، بافت و گرادیان استفاده شده است که تأکید بیشتری بر ساختارهای لبه‌ای دارد.

۳-۱-۲- معماری شبکه و فرآیند آموزش

برای یادگیری نگاشت پیچیده از تصویر ورودی به نقشه هدف G ، یک شبکه کاملاً کانولوشنی با معماری U-Net استفاده شده است. این معماری از دو بخش اصلی تشکیل شده است:

مسیر انقباضی^۱: وظیفه این مسیر، استخراج سلسله‌مراتبی ویژگی‌های بصری از تصویر ورودی است. این مسیر شامل دو بلوک اصلی است که هر یک از دو لایه کانولوشن 3×3 با تابع فعال‌سازی ReLU و لایه نرمال‌سازی دسته‌ای^۲ تشکیل شده است. تعداد فیلترها در بلوک اول ۳۲ و در بلوک دوم ۶۴ است. در انتهای هر بلوک، یک لایه Max Pooling با ابعاد 2×2 و گام ۲، ابعاد فضایی نقشه‌های ویژگی را به نصف کاهش می‌دهد. در عمیق‌ترین بخش شبکه (بخش گلوگاهی)، دو لایه کانولوشن با ۱۲۸ فیلتر بدون کاهش ابعاد، انتزاعی‌ترین ویژگی‌های تصویر را استخراج می‌کنند.

مسیر انبساطی^۳: وظیفه این مسیر، بازسازی یک نقشه با ابعاد تصویر ورودی از روی ویژگی‌های انتزاعی استخراج شده است. این کار با استفاده از لایه‌های کانولوشن انتقالی^۴ با هسته‌های 2×2 و گام ۲ انجام می‌شود که به تدریج ابعاد فضایی را افزایش می‌دهند. برای جبران جزئیات مکانی که در طی فرآیند کاهش ابعاد از دست رفته‌اند، از اتصالات پرشی^۵ استفاده می‌شود. این اتصالات، خروجی لایه‌های متناظر در مسیر انقباضی را مستقیماً به ورودی بلوک‌های مسیر انبساطی اضافه می‌کنند و به این ترتیب، شبکه می‌تواند اطلاعات دقیق مربوط به لبه‌ها و بافت‌ها را که برای تولید یک نقشه وزن باکیفیت ضروری است، حفظ کند.

لایه خروجی: در انتهای شبکه، یک لایه کانولوشن با فیلتر 1×1 قرار دارد که نقشه‌های ویژگی چندکاناله را به یک خروجی تک‌کاناله تبدیل می‌کند. این خروجی از یک لایه رگرسیون عبور داده می‌شود که تابع هزینه آن، میانگین شبکه با کمینه‌سازی تابع خطای میانگین مربعات (MSE) بین خروجی $\hat{G}$ و هدف G مطابق با رابطه (۴) آموزش می‌بیند.

$$\mathcal{L}_{MSE} = \frac{1}{N} \sum_{i=1}^N (\hat{G}_i - G_i)^2 \quad (4)$$

که در آن N تعداد کل پیکسل‌ها است. پس از آموزش، این شبکه قادر است برای هر تصویر دلخواه یک نقشه وزن W تولید کند که مقادیر آن بازتابی از مناسب بودن هر ناحیه برای جاسازی است.

فرآیند آموزش بر روی یک مجموعه داده شامل ۲۰۰۰ نمونه تصویری با ابعاد 128×128 پیکسل که از تصاویر استاندارد تولید شده‌اند، انجام می‌پذیرد. شبکه با استفاده از بهینه‌ساز آدام^۶، نرخ یادگیری اولیه ۰.۰۰۱ و یک استراتژی کاهش گام‌به‌گام نرخ یادگیری^۷ آموزش می‌بیند. بر اساس این استراتژی، نرخ یادگیری هر ۱۰ دوره با ضریب ۰.۵ کاهش می‌یابد. این مکانیزم اجازه می‌دهد که شبکه در مراحل اولیه آموزش با

¹ Encoder

² Batch Normalization

³ Decoder

⁴ Convolution Transpose

⁵ Skip Connections

⁶ Adam

⁷ Step-wise Decay

گام‌های بلند به سرعت به همگرایی نزدیک شود و در مراحل پایانی، با گام‌های کوچک‌تر به تنظیم دقیق وزن‌ها بپردازد. برای جلوگیری از بیش‌برازش، از تکنیک‌های تنظیم‌سازی L_2 و محدودسازی گرادیان¹ استفاده شده است. پس از اتمام موفقیت‌آمیز فرآیند آموزش که جزئیات و معیارهای عملکرد آن به طور کامل در بخش ۴-۲ تحلیل خواهد شد، شبکه آموزش دیده قادر است برای هر تصویر دلخواه، یک نقشه وزن تطبیقی با دقت بالا تولید کند.

۳-۲- فاز دوم: جاسازی تطبیقی بهینه‌شده در حوزه موجک با استفاده از مدل جانشین و PSO

پس از تکمیل فاز تحلیل و آماده‌سازی نقشه راهنمای جاسازی، فاز دوم که هسته اصلی عملیات نهان‌نگاری است، آغاز می‌شود. این فاز از یک ساختار ترکیبی و هوشمند بهره می‌برد که در آن، دو مؤلفه کلیدی فاکتور مقیاس بهینه و نقشه وزن تطبیقی با یکدیگر تلفیق می‌شوند تا جاسازی واترمارک به صورت کاملاً تطبیقی و مقاوم در حوزه تبدیل موجک گسسته (DWT) صورت پذیرد. رویکرد اصلی در این فاز، بهینه‌سازی غیرمستقیم و سریع پارامتر حساس سیستم با استفاده از یک مدل جانشین² است که جایگزین ارزیابی‌های مکرر و پرهزینه می‌شود.

۳-۲-۱- پیش‌پردازش و انتخاب کانال میزبان

تصویر میزبان رنگی (I_{host}) در فضای رنگی استاندارد RGB دریافت می‌شود. اولین گام در فرآیند جاسازی، انتخاب یک کانال رنگی مناسب است. بر اساس مطالعات گسترده در حوزه HVS، چشم انسان کمترین حساسیت را به تغییرات در کانال آبی دارد. بنابراین، برای بیشینه‌سازی شفافیت و نامحسوس بودن تغییرات، کانال آبی (B_{host}) به عنوان بستر اصلی جاسازی انتخاب می‌شود. این انتخاب، یک لایه تضمین اولیه برای حفظ کیفیت بصری تصویر ایجاد می‌کند.

۳-۲-۲- تجزیه موجک و انتخاب زیرباند مقاوم

بر روی کانال آبی انتخاب‌شده B_{host} ، یک تبدیل موجک گسسته دوبعدی دو سطحی با موجک Haar اعمال می‌شود. این تبدیل، تصویر را به یک مجموعه از زیرباندهای تقریبی (LL) و جزئیات LL_2, LH_2, HL_2, HH_2 مطابق با رابطه (۵) در سطوح مختلف تجزیه می‌کند. از میان این زیرباندها، زیرباند LL_2 به عنوان میزبان نهایی برای جاسازی انتخاب می‌شود. انتخاب این زیرباند مبتنی بر چند دلیل راهبردی است: LL_2 یک نسخه فشرده و پایین‌گذر از تصویر است که بخش عمده انرژی و اطلاعات ساختاری تصویر را در خود متمرکز کرده است. این ویژگی، مقاومت ذاتی بالایی در برابر بسیاری از حملات پردازش سیگنال، به ویژه فشرده‌سازی JPEG و فیلترینگ پایین‌گذر، فراهم می‌آورد.

$$[LL_2, LH_2, HL_2, HH_2] = DWT_2(B_{host}) \quad (5)$$

۳-۲-۳- چارچوب بهینه‌سازی مبتنی بر مدل جانشین

تعیین مقدار بهینه فاکتور مقیاس (α)، پارامتری که قدرت جاسازی را کنترل می‌کند، یک مسئله بهینه‌سازی پیوسته و چندهدفه است. یک α بسیار کوچک، مقاومت را قربانی شفافیت می‌کند و یک α بسیار بزرگ، کیفیت بصری را تخریب می‌کند. حل این مسئله با روش‌های سنتی سعی و خطا یا حتی با الگوریتم‌های فراابتکاری استاندارد، به دلیل نیاز به ارزیابی‌های مکرر و پرهزینه، عملاً ناکارآمد است. هر ارزیابی نیازمند اجرای کامل فرآیند جاسازی، اعمال مجموعه‌ای از حملات، استخراج واترمارک و محاسبه معیارهای خطا است.

برای غلبه بر این مانع، یک چارچوب بهینه‌سازی دومرحله‌ای نوآورانه طراحی شده است که از یک مدل جانشین بهره می‌برد. مدل جانشین، یک مدل تقریبی و سریع است که جایگزین تابع هزینه واقعی و پرهزینه می‌شود. فرآیند کار به این صورت است:

¹ Gradient Clipping

² Surrogate Model

نمونه برداری و ساخت مجموعه داده آموزشی: در گام اول، n نمونه تصادفی از فاکتور مقیاس (α_i) در بازه جستجوی $\alpha \in [\alpha_{min}, \alpha_{max}]$ تولید می شود. برای هر یک از این نمونه ها، یک بار تابع هزینه واقعی $F(\alpha_i)$ با اجرای فرآیند کامل نهان نگاری بر روی یک زیرمجموعه کوچک از تصاویر آموزشی محاسبه می شود. تابع هزینه واقعی به گونه ای تعریف می شود که مصالحه میان شفافیت و مقاومت را منعکس کند. در اینجا، این تابع به صورت نسبت اعوجاج (MSE) به توان جاسازی (α) مطابق با رابطه (۶) مدل می شود:

$$F(\alpha) = \frac{MSE(\alpha)}{\alpha + \epsilon} \quad (6)$$

این کار، یک مجموعه داده کوچک اما ارزشمند $D = \{(\alpha_i, F_i)\}_{i=1}^n$ ایجاد می کند.

آموزش مدل جانشین: در گام دوم، یک مدل درخت رگرسیون بر روی مجموعه داده D آموزش داده می شود. درخت رگرسیون یک مدل غیرپارامتری و انعطاف پذیر است که می تواند روابط غیرخطی پیچیده را بدون نیاز به فرضیات اولیه مدل کند. مهم تر از آن، ارزیابی (پیش بینی) این مدل بسیار سریع است. مدل آموزش دیده به عنوان تخمین گر $\hat{F}(\alpha)$ عمل می کند و می تواند هزینه متناظر با هر α جدید را در کسری از ثانیه پیش بینی کند.

بهینه سازی با PSO: اکنون که یک تخمین گر سریع و کم هزینه در اختیار است، از الگوریتم PSO برای کاوش کامل فضای جستجو استفاده می شود. PSO یک الگوریتم فراابتکاری مبتنی بر جمعیت است که از رفتار اجتماعی پرندگان الهام گرفته شده است. در این چارچوب، هر ذره یک کاندیدا برای α است. در هر تکرار الگوریتم، موقعیت هر ذره با فراخوانی $\hat{F}$ ارزیابی می شود، نه تابع واقعی F . این نوآوری کلیدی باعث می شود که امکان اجرای هزاران ارزیابی در کسری از زمان مورد نیاز برای یک ارزیابی واقعی فراهم شود. به این ترتیب، PSO می تواند به طور کامل فضای جستجو را بکاود و به سمت α_{opt} همگرا شود. فرآیند به روزرسانی سرعت و موقعیت ذرات طبق روابط استاندارد PSO انجام می شود.

اکنون که یک مدل جانشین سریع $\hat{F}$ در اختیار داریم، از الگوریتم بهینه سازی ازدحام ذرات (PSO) برای یافتن α_{opt} استفاده می کنیم که $\hat{F}$ را کمینه می کند. PSO یک الگوریتم فراابتکاری مبتنی بر جمعیت است که از رفتار اجتماعی پرندگان الهام گرفته شده است. هر ذره در فضای جستجو (که در اینجا فضای یک بعدی α است) یک راه حل کاندیدا را نشان می دهد. موقعیت یک ذره با p و سرعت آن با v نمایش داده می شود.

در هر تکرار t ، سرعت و موقعیت هر ذره i با استفاده از روابط (۷) به روزرسانی می شود:

$$\begin{aligned} v_i(t+1) &= \omega \cdot v_i(t) + c_1 r_1 (pbest_i - p_i(t)) + c_2 r_2 (gbest - p_i(t)) \\ p_i(t+1) &= p_i(t) + v_i(t+1) \end{aligned} \quad (7)$$

که در آن:

- ω وزن اینرسی که نرخ اکتشاف و بهره برداری را کنترل می کند.
- c_1, c_2 ضرایب یادگیری شناختی و اجتماعی هستند.
- r_1, r_2 اعداد تصادفی با توزیع یکنواخت در بازه $[0,1]$ می باشند.
- $pbest_i$ بهترین موقعیتی است که ذره i تاکنون تجربه کرده است.
- $gbest$ بهترین موقعیتی است که در کل جمعیت یافت شده است.

در هر تکرار، هزینه یک موقعیت جدید p_i با فراخوانی تابع تخمین گر سریع $\hat{F}(p_i)$ نه تابع واقعی F ارزیابی می شود. این مهم ترین برتری روش ما است که اجازه می دهد PSO با هزاران ارزیابی، فضای جستجو را بدون هیچ هزینه محاسباتی قابل توجهی بکاهد. پس از هم گرایی الگوریتم، $gbest$ نهایی به عنوان فاکتور مقیاس بهینه α_{opt} انتخاب می شود.

• ۳-۲-۴- جاسازی تطبیقی نهایی

پس از تعیین α_{opt} توسط حلقه PSO و تولید نقشه وزن (W) توسط CNN، این دو در یک گام جاسازی یکپارچه و تطبیقی تلفیق می شوند. برای این کار، ابتدا نقشه وزن W (که در فاز اول با ابعاد 128×128 تولید شده) به ابعاد زیرباند $LL2LL2$ نمونه برداری مجدد (Resize) می شود تا یک تناظر یک به یک بین وزن ها و ضرایب موجک برقرار گردد. همچنین، واترمارک باینری (WM_{binary}) که پیش تر با نگاشت آرنولد¹ رمزنگاری و درهم سازی شده، به همان ابعاد تنظیم می شود. عملیات جاسازی با استفاده از رابطه (۸) انجام می پذیرد:

$$LL'_2(x, y) = LL_2(x, y) + (\alpha_{opt} \cdot W_{norm}(x, y)) \cdot WM_{binary}(x, y) \quad (8)$$

این معادله، نقطه اوج روش پیشنهادی و تجلی بخش ماهیت تطبیقی آن است. قدرت جاسازی مؤثر در هر ضریب موجک، حاصل ضرب دو عامل کاملاً متمایز است:

عامل سراسری (α_{opt}): یک مقدار اسکالر که توسط فرآیند بهینه سازی مبتنی بر مدل جانشین و PSO برای کل تصویر به صورت بهینه تعیین شده است و تضمین کننده تعادل کلی بین شفافیت و مقاومت است.

عامل محلی ($W_{norm}(x, y)$): یک وزن تطبیقی که توسط شبکه CNN برای هر پیکسل مشخص شده است و میزان پیچیدگی بصری و ظرفیت پنهان سازی آن ناحیه خاص را منعکس می کند.

این مکانیزم هوشمندانه تضمین می کند که در نواحی صاف و حساس تصویر که W_{norm} به صفر نزدیک است، قدرت جاسازی حداقل باشد و کیفیت بصری حفظ شود. در مقابل، در نواحی پر جزئیات، لبه ها و بافت های پیچیده که W_{norm} به یک نزدیک است، بخش اعظم انرژی واترمارک به صورت مقاوم جاسازی گردد.

• ۳-۲-۵- بازسازی تصویر نهان نگاری شده

در گام نهایی، زیرباند تقریبی اصلاح شده (LL'_2) جایگزین زیرباند اصلی (LL_2) در بردار کامل ضرایب موجک شده و سپس تبدیل موجک معکوس دوبعدی (IDWT) بر روی آن اعمال می گردد تا کانال آبی جدید (B'_{host}) مطابق با رابطه (۹) بازسازی شود.

$$B'_{host} = IDWT_2(LL'_2, LH_2, HL_2, HH_2) \quad (9)$$

در نهایت، این کانال جدید با کانال های قرمز و سبز دست نخورده تصویر اصلی ترکیب شده و تصویر نهان نگاری شده نهایی (I_{WM}) را مطابق با رابطه (۱۰) تشکیل می دهد. این تصویر باید به طور هم زمان دو شرط اساسی را برآورده سازد: از نظر بصری از تصویر اصلی غیر قابل تمایز باشد (شفافیت) و واترمارک جاسازی شده در آن بتواند در برابر حملات مختلف، به ویژه حملات هندسی، باز یابی شود (مقاومت). ارزیابی دقیق این ویژگی ها موضوع اصلی بخش ۴ خواهد بود.

$$I_{WM} = merge(R_{host}, G_{host}, B'_{host}) \quad (10)$$

¹ Arnold Transform

۴- یافته‌های پژوهش: نتایج تجربی، تحلیل و مقایسه

برای تأیید کارایی و برتری روش پیشنهادی، یک سری آزمایش‌های جامع طراحی و اجرا شده است. این بخش به تشریح تنظیمات آزمایش، ارائه نتایج کمی و کیفی و تحلیل عمیق آن‌ها می‌پردازد.

۴-۱- تنظیمات آزمایش و معیارهای ارزیابی

۴-۱-۱- مجموعه داده و پیش‌پردازش

برای آموزش تحلیل‌گر صحنه هوشمند، یک مجموعه داده اختصاصی با بهره‌گیری از تصاویر استاندارد پرکاربرد در حوزه پردازش تصویر تولید شد. تصاویر میزبان اصلی شامل شش تصویر رنگی 512×512 پیکسلی استاندارد با عناوین Jet, Boat, Barbara, Baboon, Lena و Peppers به همراه یک تصویر لوگو (مورد استفاده در فاز دوم به عنوان واترمارک) می‌باشند. برای افزایش تنوع و غنای مجموعه داده آموزشی و جلوگیری از جلوگیری از بیش‌برازش، فرآیند افزایش داده به صورت آفلاین انجام پذیرفت. بدین منظور، ۲۰۰۰ نمونه تصویری با ابعاد 128×128 پیکسل از طریق برش‌های تصادفی، چرخش‌های جزئی و بازتاب‌های افقی/عمودی از تصاویر اصلی استخراج گردید. برای هر یک از این نمونه‌ها، یک نقشه حقیقت مبنا که بیانگر شایستگی هر پیکسل برای جاسازی مقاوم و شفاف است، تولید شد. این فرآیند با رویکردی ترکیبی و مبتنی بر اصول HVS صورت پذیرفت. منطق اساسی این است که نواحی با پیچیدگی بصری بالا (نظیر لبه‌ها و بافت‌ها) به دلیل اثر پوشاندگی، مستعدترین مناطق برای جاسازی هستند. بر این اساس، نقشه هدف از ترکیب خطی سه معیار بینایی محاسباتی زیر ساخته شد:

نقشه لبه: با استفاده از آشکارساز لبه کنی با آستانه‌های $[0.05, 0.15]$ ، نقشه باینری لبه‌های قوی تصویر استخراج گردید (وزن ۰.۴).
نقشه بافت: با اعمال فیلتر انحراف معیار محلی^۱ با پنجره‌ای به ابعاد 7×7 ، میزان تغییرات شدت روشنایی به عنوان معیاری از پیچیدگی بافت اندازه‌گیری و نرمال‌سازی شد (وزن ۰.۳).
نقشه بزرگی گرادیان: بزرگی بردار گرادیان تصویر با استفاده از عملگر سوبل محاسبه گردید تا بر قدرت لبه‌ها تأکید بیشتری شود (وزن ۰.۳).

نقشه ترکیبی نهایی با یک فیلتر گاوسی با انحراف معیار ۲ پیکسل هموار و سپس به بازه $[0, 1]$ نرمال‌سازی شد. مجموعه داده نهایی به صورت تصادفی به دو بخش آموزش (۸۵٪، معادل ۱۷۰۰ نمونه) و اعتبارسنجی (۱۵٪، معادل ۳۰۰ نمونه) تقسیم گردید.

۴-۱-۲- معماری و تنظیمات شبکه عصبی کانولوشنی

همان‌طور که پیش‌تر اشاره شد، برای تحلیل هوشمند تصویر میزبان از یک شبکه کاملاً کانولوشنی با معماری U-Net استفاده شده است. جزئیات کامل این معماری در جداول ۱ تا ۳ ارائه گردیده است. جدول ۱، ساختار مسیر انقباضی را نشان می‌دهد که وظیفه استخراج سلسله‌مراتبی ویژگی‌های بصری را بر عهده دارد. جدول ۲، جزئیات بخش میانی و مسیر انبساطی را که به بازسازی نقشه وزن با وضوح کامل می‌پردازد، مشخص می‌کند. در نهایت، جدول ۳ خلاصه‌ای از پارامترهای کلیدی آموزش شبکه را ارائه می‌دهد.

¹ Local Standard Deviation

جدول ۱ جزئیات لایه‌های مسیر انقباضی (Encoder) معماری U-Net پیشنهادی

شماره لایه	نام لایه	نوع لایه	ابعاد خروجی	تعداد فیلتر / اندازه کرنل / استراید
۱	input	Image Input	$128 \times 128 \times 3$	-
۲	enc_conv1	Convolution	$128 \times 128 \times 32$	۳۲ فیلتر / $3 \times 3 \times 3$ / استراید [۱ ۱]
۳	enc_bn1	Batch Normalization	$128 \times 128 \times 32$	۳۲ کانال
۴	enc_relu1	ReLU	$128 \times 128 \times 32$	-
۵	enc_conv2	Convolution	$128 \times 128 \times 32$	۳۲ فیلتر / $3 \times 3 \times 3$ / استراید [۱ ۱]
۶	enc_bn2	Batch Normalization	$128 \times 128 \times 32$	۳۲ کانال
۷	enc_relu2	ReLU	$128 \times 128 \times 32$	-
۸	enc_pool1	Max Pooling	$64 \times 64 \times 32$	۲×۲ استراید [۲ ۲]
۹	enc_conv3	Convolution	$64 \times 64 \times 64$	۶۴ فیلتر / $3 \times 3 \times 3$ / استراید [۱ ۱]
۱۰	enc_bn3	Batch Normalization	$64 \times 64 \times 64$	۶۴ کانال
۱۱	enc_relu3	ReLU	$64 \times 64 \times 64$	-
۱۲	enc_conv4	Convolution	$64 \times 64 \times 64$	۶۴ فیلتر / $3 \times 3 \times 3$ / استراید [۱ ۱]
۱۳	enc_bn4	Batch Normalization	$64 \times 64 \times 64$	۶۴ کانال
۱۴	enc_relu4	ReLU	$64 \times 64 \times 64$	-
۱۵	enc_pool2	Max Pooling	$32 \times 32 \times 64$	۲×۲ استراید [۲ ۲]

جدول ۲ جزئیات لایه‌های بخش میانی (Bottleneck) و مسیر انبساطی (Decoder) معماری U-Net پیشنهادی

شماره لایه	نام لایه	نوع لایه	ابعاد خروجی (Activations)	تعداد فیلتر / اندازه کرنل / استراید
۱۶	bottleneck_conv1	Convolution	$32 \times 32 \times 128$	۱۲۸ فیلتر / $64 \times 3 \times 3$ / استراید [۱ ۱]
۱۷	bottleneck_bn1	Batch Normalization	$32 \times 32 \times 128$	۱۲۸ کانال
۱۸	bottleneck_relu1	ReLU	$32 \times 32 \times 128$	-
۱۹	bottleneck_conv2	Convolution	$32 \times 32 \times 128$	۱۲۸ فیلتر / $128 \times 3 \times 3$ / استراید [۱ ۱]
۲۰	bottleneck_bn2	Batch Normalization	$32 \times 32 \times 128$	۱۲۸ کانال
۲۱	bottleneck_relu2	ReLU	$32 \times 32 \times 128$	-
۲۲	dec_upconv1	Transposed Conv.	$64 \times 64 \times 64$	۶۴ فیلتر / $128 \times 2 \times 2$ / استراید [۲ ۲]
۲۳	dec_bn1	Batch Normalization	$64 \times 64 \times 64$	۶۴ کانال
۲۴	dec_relu1	ReLU	$64 \times 64 \times 64$	-
۲۵	dec_conv1	Convolution	$64 \times 64 \times 64$	۶۴ فیلتر / $64 \times 3 \times 3$ / استراید [۱ ۱]
۲۶	dec_bn2	Batch Normalization	$64 \times 64 \times 64$	۶۴ کانال
۲۷	dec_relu2	ReLU	$64 \times 64 \times 64$	-
۲۸	dec_upconv2	Transposed Conv.	$128 \times 128 \times 32$	۳۲ فیلتر / $64 \times 2 \times 2$ / استراید [۲ ۲]
۲۹	dec_bn3	Batch Normalization	$128 \times 128 \times 32$	۳۲ کانال
۳۰	dec_relu3	ReLU	$128 \times 128 \times 32$	-
۳۱	dec_conv2	Convolution	$128 \times 128 \times 32$	۳۲ فیلتر / $32 \times 3 \times 3$ / استراید [۱ ۱]
۳۲	dec_bn4	Batch Normalization	$128 \times 128 \times 32$	۳۲ کانال
۳۳	dec_relu4	ReLU	$128 \times 128 \times 32$	-
۳۴	output_conv	Convolution	$128 \times 128 \times 1$	۱ فیلتر / $32 \times 1 \times 1$ / استراید [۱ ۱]
۳۵	output	Regression (MSE)	$128 \times 128 \times 1$	-

خلاصه ای از فرآپارامترهای اصلی استفاده شده برای آموزش شبکه در جدول ۳ ارائه شده است. این پارامترها بر اساس آزمایش های اولیه و با هدف دستیابی به همگرایی پایدار و دقت بالا انتخاب شده اند.

جدول ۳ خلاصه فرآپارامترهای اصلی آموزش شبکه U-Net.

پارامتر	مقدار
بهینه ساز	Adam
نرخ یادگیری اولیه	۰.۰۰۱
استراتژی کاهش نرخ یادگیری	(piecewise کاهش گام به گام)
ضریب کاهش نرخ یادگیری	۰.۵
دوره کاهش نرخ یادگیری	هر ۱۰ Epoch
حداکثر تعداد Epoch	۲۰
اندازه Mini-Batch	۱۶
تابع هزینه	Mean Squared Error (MSE)
تنظیم سازی L2	۰.۰۰۰۱
محدود سازی گرادین	۱.۰

همان طور که در جداول ۱ و ۲ مشاهده می شود، معماری شبکه مقارن بوده و در مجموع شامل ۳۵ لایه است. استفاده از اتصالات پرشی (که در کد پیاده سازی شده، اما برای سادگی در جداول نیامده است) بین مسیرهای انقباضی و انبساطی، به شبکه اجازه می دهد تا جزئیات مکانی از دست رفته در فرآیند کاهش ابعاد را بازیابی کند. این طراحی، شبکه را قادر می سازد تا نگاشت پیچیده از تصویر میزبان به نقشه وزن تطبیقی را با دقت بالا (میانگین ۹۹.۹۸۶٪ در مجموعه اعتبارسنجی) یاد بگیرد، که مبنای عملکرد موفق فاز دوم سیستم نهان نگاری را فراهم می آورد.

۴-۱-۳- معیارهای ارزیابی

برای ارزیابی همه جانبه و کمی عملکرد چارچوب نهان نگاری پیشنهادی، از دو دسته معیار استاندارد استفاده شده است که به ترتیب به سنجش شفافیت (کیفیت بصری تصویر نهان نگاری شده) و مقاومت (صحت بازیابی واترمارک پس از اعمال حملات) می پردازند. این معیارها امکان مقایسه دقیق روش پیشنهادی با پژوهش های پیشین را فراهم می آورند.

الف) معیارهای سنجش شفافیت: هدف این دسته از معیارها، کمی سازی میزان اعوجاج تحمیل شده بر تصویر میزبان در نتیجه فرآیند جاسازی واترمارک است. یک سیستم نهان نگاری ایده آل باید ضمن بیشینه سازی مقاومت، کمترین تخریب ممکن را در کیفیت بصری تصویر ایجاد کند.

نسبت سیگنال به نویز بیشینه (PSNR): این معیار، که یکی از رایج ترین سنجها در حوزه پردازش تصویر است، نسبت بین حداکثر توان ممکن یک سیگنال (تصویر اصلی) به توان نویز مخرب (اعوجاج ناشی از جاسازی) را بر حسب دسی بل (dB) اندازه گیری می کند. مقادیر بالاتر PSNR به طور معمول بیانگر شباهت بیشتر تصویر نهان نگاری شده به تصویر اصلی و در نتیجه، شفافیت بهتر است. در کاربردهای نهان نگاری، PSNR بالای ۳۸ دسی بل نشان دهنده نامحسوس بودن تغییرات برای سیستم بینایی انسان تلقی می شود. این معیار مستقیماً از معیار MSE مشتق می شود و رابطه آن به صورت معادله (۱۱) تعریف می گردد:

$$PSNR = 10 \log_{10} \left(\frac{MAX_I^2}{MSE} \right) \quad (11)$$

که در آن، MAX_I^2 حداکثر مقدار شدت روشنایی پیکسل‌های تصویر (برای یک تصویر ۸-بیتی برابر با ۲۵۵) و MSE میانگین مربعات خطا است که در ادامه تعریف می‌شود.

میانگین مربعات خطا (MSE): این معیار، میانگین مجذور اختلاف میان یکایک پیکسل‌های تصویر اصلی I_{host} و تصویر نهان‌نگاری شده I_{WM} را محاسبه می‌کند. MSE مبنای محاسبه PSNR بوده و به طور مستقیم میزان انرژی اعوجاج را کمی‌سازی می‌نماید. مقادیر کمتر MSE بیانگر اعوجاج کمتر و کیفیت بالاتر است. فرمول محاسباتی آن در معادله (۱۲) ارائه شده است:

$$MSE = \frac{1}{M * N} \sum_{i=1}^M \sum_{j=1}^N [I_{host}(i, j) - I_{WM}(i, j)]^2 \quad (12)$$

که در آن، M و N به ترتیب تعداد سطرها و ستون‌های تصویر را نشان می‌دهند.

شاخص تشابه ساختاری (SSIM): برخلاف PSNR و MSE که صرفاً بر اساس خطای مطلق پیکسل‌ها عمل می‌کنند، SSIM یک معیار ادراکی است که تطابق میان تصویر اصلی و تصویر نهان‌نگاری شده را از سه منظر «درخشندگی»، «کنتراست» و «ساختار» ارزیابی می‌کند. این شاخص با الهام از نحوه درک سیستم بینایی انسان از اطلاعات ساختاری صحنه طراحی شده و مقدار آن در بازه $[0, 1]$ قرار دارد. مقدار ۱ نشان‌دهنده تشابه ساختاری کامل است و آن را به معیاری قوی‌تر برای سنجش شفافیت تبدیل می‌کند.

ب) معیارهای سنجش مقاومت: این دسته از معیارها، توانایی الگوریتم در بازیابی صحیح واترمارک جاسازی شده را پس از آنکه تصویر نهان‌نگاری شده تحت حملات عمدی یا غیرعمدی (مانند چرخش و فشردگی) قرار گرفت، ارزیابی می‌کنند. هسته اصلی این ارزیابی، مقایسه واترمارک اصلی (WM) با واترمارک استخراج شده (WM') است.

ضریب همبستگی نرمال شده (NC): این معیار، میزان شباهت و همبستگی خطی میان واترمارک اصلی و واترمارک استخراج شده را اندازه می‌گیرد. مقدار NC از ۰ (عدم تشابه) تا ۱ (تشابه کامل) متغیر است. در یک سیستم مقاوم، حتی پس از اعمال حملات مخرب، انتظار می‌رود مقدار NC بسیار نزدیک به ۱ (یا ۱۰۰٪) باقی بماند که نشان‌دهنده صحت بالای فرآیند استخراج است. این معیار با استفاده از رابطه (۱۳) محاسبه می‌شود:

$$NC = \frac{\sum_{i=1}^P \sum_{j=1}^Q [WM(i, j) \cdot WM'(i, j)]}{\sqrt{\sum_{i=1}^P \sum_{j=1}^Q [WM(i, j)^2]} \sqrt{\sum_{i=1}^P \sum_{j=1}^Q [WM'(i, j)^2]}} \quad (13)$$

که در آن، P و Q ابعاد واترمارک هستند.

نرخ خطای بیت (BER): این معیار، نسبت تعداد بیت‌هایی از واترمارک که به اشتباه استخراج شده‌اند به کل بیت‌های واترمارک را نشان می‌دهد. BER یک معیار مستقیم از خطای بازیابی است و با NC رابطه معکوس دارد؛ به طوری که یک سیستم نهان‌نگاری ایده‌آل باید BER را در کمترین حد ممکن (نزدیک به صفر) نگه دارد. نحوه محاسبه آن در معادله (۱۴) آمده است:

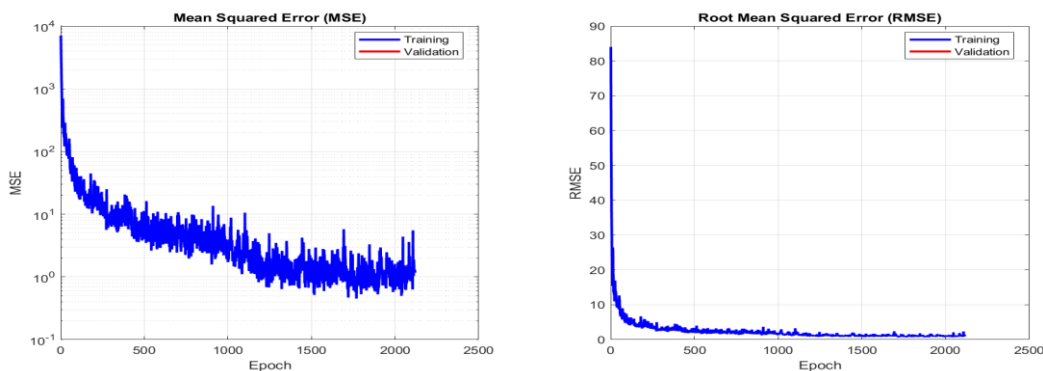
$$BER = \frac{1}{P * Q} \sum_{i=1}^P \sum_{j=1}^Q [WM(i, j) \oplus WM'(i, j)] \quad (14)$$

که در آن، $\oplus$ عملگر منطقی XOR را نشان می‌دهد.

۴-۲- تحلیل نتایج فاز اول: آموزش شبکه عصبی و تولید نقشه‌های وزن تطبیقی

در این بخش، به ارزیابی دقیق عملکرد فاز اول روش پیشنهادی که شامل آموزش شبکه عصبی کانولوشنی عمیق برای تولید نقشه‌های وزن تطبیقی است، پرداخته می‌شود. تمامی تحلیل‌های ارائه شده در این زیربخش، مربوط به فرآیند آموزش و ارزیابی شبکه CNN پیش از ورود به فاز دوم نهان‌نگاری می‌باشد. شبکه مذکور بر روی مجموعه داده‌ای شامل ۲۰۰۰ نمونه (۱۷۰۰ نمونه آموزش و ۳۰۰ نمونه اعتبارسنجی) با ابعاد 128×128 پیکسل و به مدت ۲۰ دوره آموزش دیده است. نتایج حاصل در قالب نمودارها و جداول متعددی گردآوری شده‌اند که در ادامه به تحلیل هر یک پرداخته می‌شود.

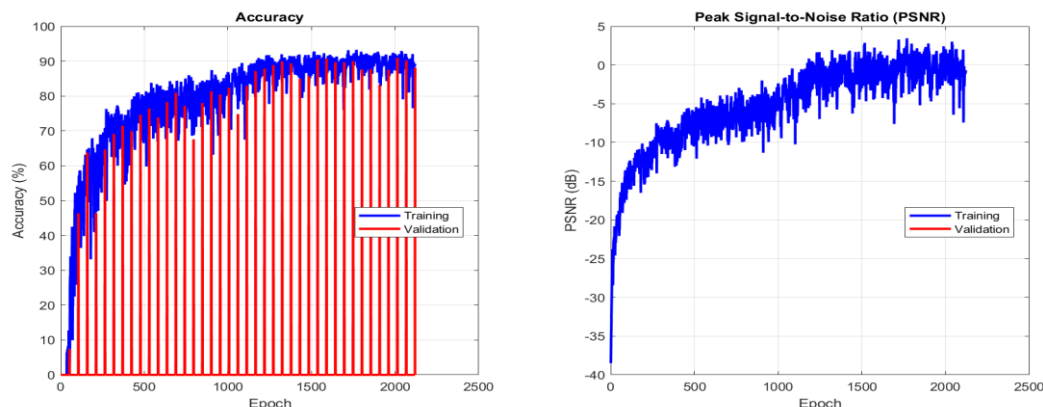
به منظور بررسی دقیق‌تر عملکرد مدل در کاهش خطای پیش‌بینی، شکل ۲ روند تغییرات دو معیار مهم MSE و RMSE را در طی ۲۰۰۰ دوره آموزش نشان می‌دهد.



شکل ۲ نمودارهای خطا

در نمودار ۲، سمت چپ شکل خطاهای آموزش، مقایسه MSE آموزش (منحنی آبی) و اعتبارسنجی (منحنی قرمز) در مقیاس لگاریتمی نشان داده شده است. نوسانات نسبتاً زیاد در دوره‌های اولیه، حاکی از یادگیری سریع و پرشتاب مدل می‌باشد. پس از گذشت حدود ۲۰۰ دوره، هر دو منحنی به طور پایدار به زیر سطح 4×10^{-1} می‌رسند. نکته قابل توجه، نزدیکی بسیار زیاد دو منحنی آموزش و اعتبارسنجی در سراسر فرآیند است که نشان‌دهنده عدم وقوع پدیده بیش‌برازش می‌باشد. نمودار سمت راست همین شکل، الگوی مشابهی را برای RMSE نمایش می‌دهد. مقادیر نهایی RMSE کمتر از ۰.۰۱ بوده و بر اساس داده‌های فایل اکسل، میانگین نهایی RMSE حدود ۰.۰۱۰۸ و میانگین MSE حدود 1.44×10^{-4} است.

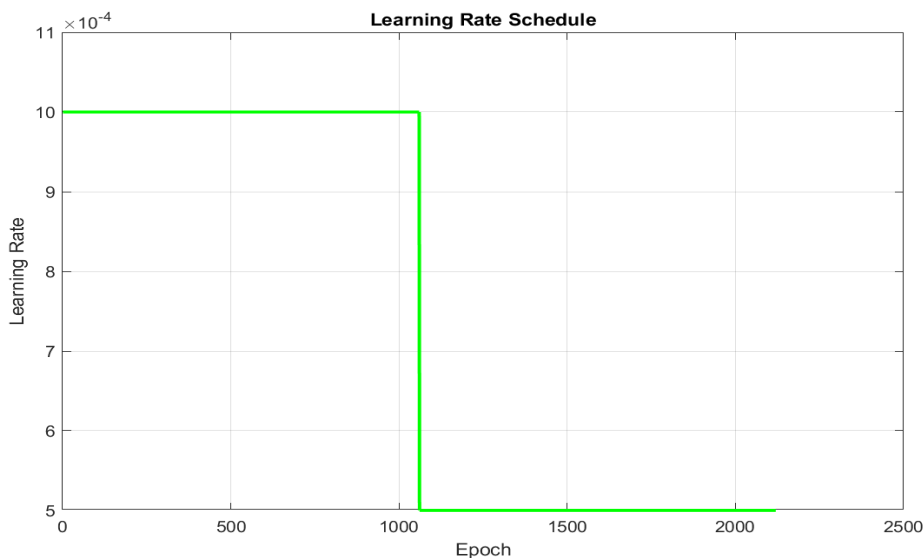
علاوه بر معیارهای خطا، ارزیابی دقت و کیفیت بازسازی نقشه‌های وزن از اهمیت ویژه‌ای برخوردار است. شکل ۳ نمودارهای دقت (بر حسب درصد) و نسبت سیگنال به نویز بیشینه (PSNR) را برای داده‌های آموزش و اعتبارسنجی به تصویر کشیده است.



شکل ۳ معیارهای عملکرد

در شکل ۳، نمودار سمت چپ به نمایش دقت آموزش و اعتبارسنجی اختصاص دارد. هر دو منحنی با شیب تند افزایش یافته و پس از حدود ۱۵۰۰ دوره، به دقت نهایی نزدیک ۱۰۰ درصد رسیده اند. طبق آمار فایل اکسل، دقت اعتبارسنجی نهایی برابر ۹۹.۹۸۶ درصد است. نمودار سمت راست نیز PSNR را نشان می دهد که به عنوان تابعی از لگاریتم معکوس MSE، رشد پایدار و یکنواختی داشته و در پایان آموزش به حدود ۴۰ دسی بل رسیده است. آمار دقیق حاکی از آن است که میانگین نهایی PSNR برابر ۴۰.۰۸ دسی بل با انحراف معیار ۳.۳۸ دسی بل می باشد. این نمودارها موفقیت فرآیند آموزش فاز اول را به طور قطعی تأیید می کنند.

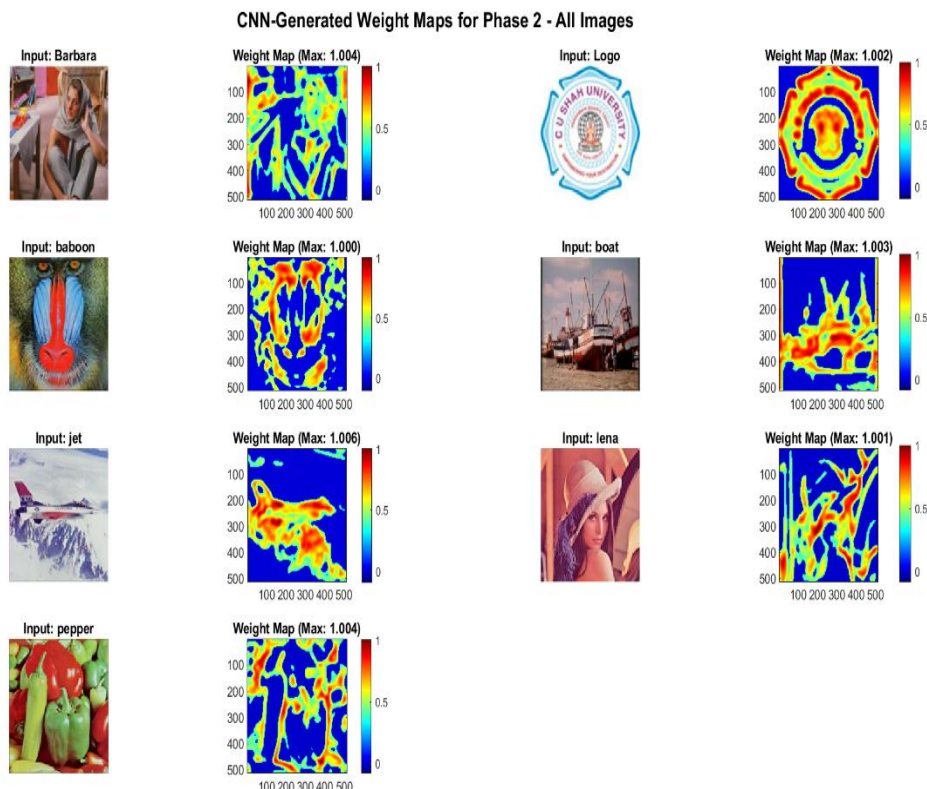
تنظیم ابرپارامتر نرخ یادگیری نقش کلیدی در همگرایی مدل دارد. شکل ۴ استراتژی کاهش گام به گام این نرخ را در طول فرآیند آموزش نمایش می دهد.



شکل ۴ استراتژی کاهش نرخ یادگیری

در نمودار ۴، محور افقی بیانگر دوره های آموزش و محور عمودی نشان دهنده مقدار نرخ یادگیری است. مشاهده می شود که نرخ یادگیری با مقدار اولیه ۰.۰۰۱ شروع شده و تا پایان دوره ۱۰۰۰ ثابت باقی می ماند. در این نقطه، نرخ یادگیری به یکباره به مقدار ۰.۰۰۰۵ کاهش می یابد (یک گام کاهشی) و تا پایان ۲۰۰۰ دوره آموزش در همین مقدار ثابت می ماند. این استراتژی رایج که با عنوان کاهش گام های نرخ یادگیری^۱ شناخته می شود، امکان همگرایی سریع در مراحل اولیه (با گام های بزرگ) و سپس تنظیم دقیق و ظریف پارامترها در مراحل پایانی (با گام های کوچک) را فراهم می آورد. انتخاب نقطه کاهش در دوره ۱۰۰۰ نیز مبتنی بر مشاهده روند همگرایی منحنی خطا بوده است. به منظور تأیید کیفی توانایی شبکه در شناسایی نواحی مناسب برای جاسازی، شکل ۵ نقشه های وزن تطبیقی تولید شده توسط CNN آموزش دیده را برای هفت تصویر میزبان اصلی نشان می دهد.

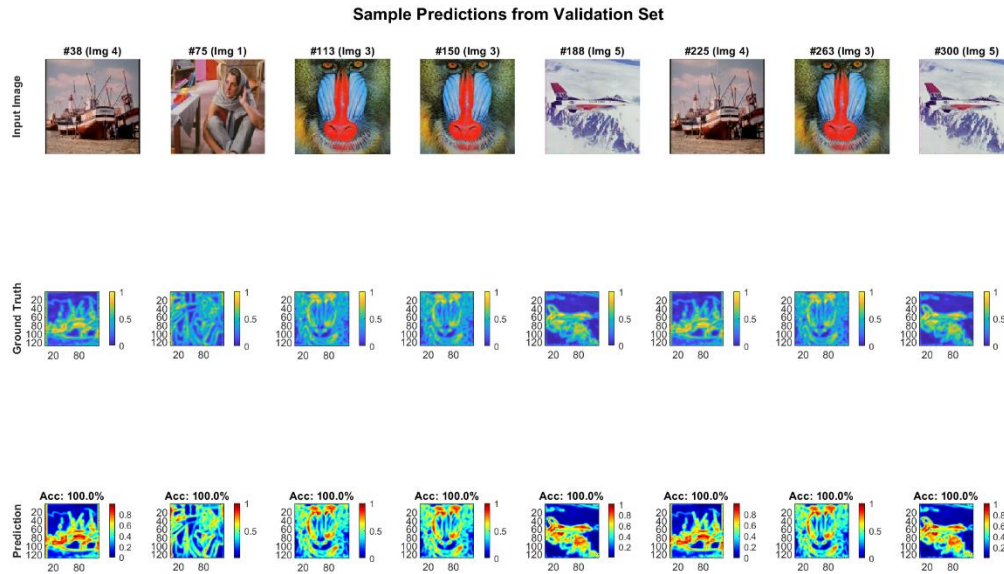
¹ Step-wise Learning Rate Decay



شکل ۵ نقشه های وزن تولید شده برای تصاویر میزبان فاز دوم

پس از اتمام فرآیند آموزش، شبکه CNN آموزش دیده برای تولید نقشه های وزن تطبیقی برای هفت تصویر میزبان مورد استفاده در فاز دوم (شامل Barbara, Baboon, Jet, Pepper, Logo, Boat و Lena) به کار گرفته شده است. نتایج در قالب شکلی ترکیبی ارائه شده که هر ردیف آن شامل دو تصویر است: تصویر ورودی در سمت چپ و نقشه وزن متناظر تولید شده توسط شبکه در سمت راست. نقشه های وزن با استفاده از طیف رنگی jet (مقادیر پایین به رنگ آبی و مقادیر بالا به رنگ قرمز) نمایش داده شده اند. بررسی این نقشه ها نشان می دهد که رنگ های گرم (نارنجی و قرمز) عمدتاً بر روی نواحی با پیچیدگی بصری بالا مانند لبه ها و بافت های شلوغ متمرکز شده اند، در حالی که رنگ های سرد (آبی) نواحی صاف و یکنواخت را پوشش می دهند. به عنوان نمونه، در تصویر Jet لبه های تیز هواپیما با وزن بالایی مشخص شده اند در حالی که ناحیه آسمان وزن نزدیک به صفر دارد. در تصاویر Barbara و Baboon، بافت های پیچیده مو و صورت با وزن بالا شناسایی شده اند. در تصویر Lena نیز لبه های کلاه، پرده ها و شانه وزن بالایی دارند. نقشه وزن مربوط به تصویر Logo نیز جزئیات لبه های لوگو را به خوبی آشکار سازی کرده است. بر اساس آمار مندرج در فایل اکسل، حداکثر وزن در این نقشه ها بین ۱.۰۰۰ تا ۱.۰۰۶ و میانگین وزن کلی حدود ۰.۳ است که تأیید می کند وزن به صورت هوشمندانه در نواحی خاص متمرکز شده و در کل تصویر توزیع یکنواختی ندارد.

برای ارزیابی توانایی تعمیم پذیری مدل به تصاویر دیده نشده، شکل ۶، نتایج پیش بینی شبکه را برای هشت نمونه تصادفی از مجموعه اعتبارسنجی، در مقایسه با نقشه های حقیقت مبنا ارائه می کند.



شکل ۶ نمونه پیش‌بینی‌های شبکه بر روی داده‌های اعتبارسنجی

به منظور ارزیابی کیفی عملکرد شبکه، هشت نمونه تصادفی از مجموعه اعتبارسنجی انتخاب شده و نتایج پیش‌بینی در قالب شکلی سه ردیفی نمایش داده شده است. ردیف اول شامل تصاویر ورودی، ردیف دوم شامل نقشه‌های حقیقت مبنا و ردیف سوم شامل نقشه‌های پیش‌بینی‌شده توسط شبکه می‌باشد. نقشه‌های ردیف دوم و سوم با استفاده از طیف رنگی jet نمایش داده شده‌اند. مقایسه مستقیم هر ستون از این شکل، شباهت بسیار بالای نقشه‌های تولیدشده توسط شبکه به نقشه‌های حقیقت مبنا را به وضوح نشان می‌دهد. دقت پیش‌بینی هر نمونه در عنوان زیر تصویر مرتبط درج شده است که برای اکثر نمونه‌ها مقداری در حدود 99.99% را نشان می‌دهد. برای مثال، در نمونه اول که تصویری با لبه‌های مشخص و برجسته است، شبکه توانسته همان الگوی لبه‌ها را با دقت بسیار بالا بازسازی کند. این نتایج به طور قطعی توانایی تعمیم‌پذیری عالی شبکه را به تصویری که در فرآیند آموزش دیده نشده‌اند، تأیید می‌کند و نشان می‌دهد که شبکه صرفاً به حافظه‌سازی تصاویر آموزش نپرداخته است.

خلاصه نتایج نهایی فاز اول، ارقام دقیق و جامعی از عملکرد مدل را ارائه می‌دهد. از نظر زمانی، کل اجرای فاز اول حدود ۱۲۰.۶۷ دقیقه (معادل ۲ ساعت) به طول انجامیده که از این مقدار، ۱۱۴.۷۶ دقیقه به خالص فرآیند آموزش اختصاص داشته است. از نظر دقت، میانگین دقت ۹۹.۹۸۶ درصد با انحراف معیار ۰.۰۱۶ درصد، کمینه دقت ۹۹.۹۴۴ درصد و بیشینه دقت ۹۹.۹۹۴ درصد به دست آمده است. این ارقام نشان‌دهنده عملکرد استثنایی، پایدار و کم‌نوسان مدل بر روی تمامی داده‌ها هستند. از نظر معیارهای خطا، میانگین MSE برابر 1.44×10^{-1} ، میانگین RMSE برابر ۰.۰۱۰۸ و میانگین MAE برابر ۰.۰۰۷۸ است که همگی اعداد بسیار کوچکی بوده و گویای دقت بالای بازسازی می‌باشند. از نظر کیفیت بازسازی نیز، میانگین PSNR برابر ۴۰.۰۸ دسی‌بل با انحراف معیار ۳.۳۸ دسی‌بل و میانگین SSIM برابر ۰.۹۹۰ با انحراف معیار ۰.۰۰۷ به دست آمده است. این اعداد به ویژه SSIM نزدیک به ۱، نشان‌دهنده حفظ تقریباً کامل ساختار تصاویر بازسازی‌شده است.

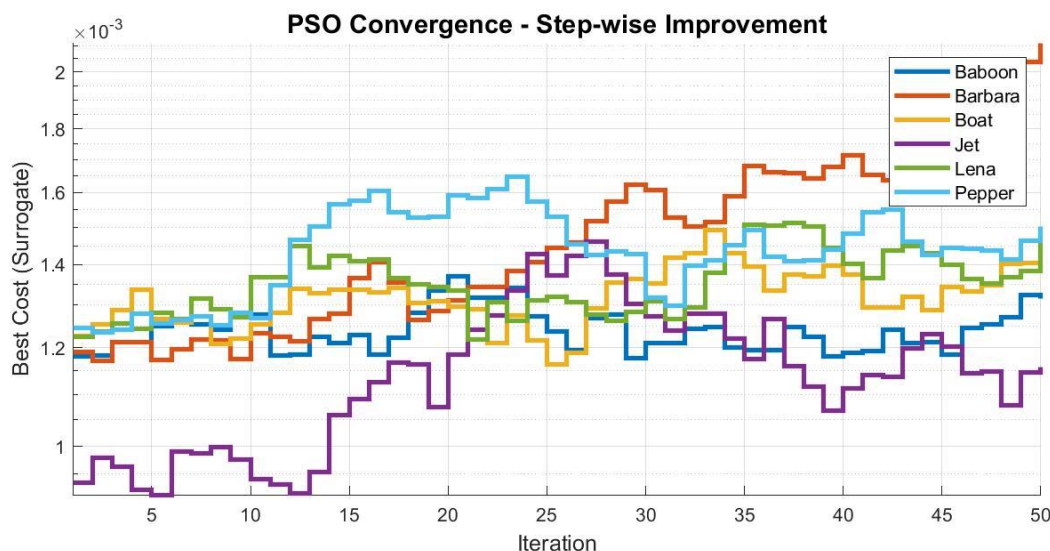
تمامی اشکال، نمودارها و جداول ارائه‌شده در این زیربخش، به طور یکپارچه و قاطع نشان می‌دهند که فاز اول روش پیشنهادی با موفقیت کامل اجرا شده است. شبکه عصبی کانولوشنی عمیق با معماری U-Net توانسته است با دقت استثنایی (میانگین ۹۹.۹۸۶ درصد)،

خطای بسیار پایین (MSE) حدود $1.4 \times 10^{-4} \sim 4$ (و همگرایی پایدار و عاری از بیش‌برازش، نگاشت پیچیده بین تصاویر ورودی و نقشه‌های وزن شایستگی را یاد بگیرد. نقشه‌های وزن تولیدشده برای تصاویر میزبان فاز دوم، به وضوح توانایی شبکه در شناسایی هوشمندانه نواحی مناسب برای جاسازی (لبه‌ها و بافت‌های پیچیده) را تأیید می‌کنند. این نقشه‌ها به عنوان خروجی نهایی فاز اول، آماده استفاده در فاز دوم روش پیشنهادی برای هدایت فرآیند جاسازی تطبیقی واترمارک در حوزه موجک خواهند بود.

۳-۴ تحلیل نتایج فاز دوم: ارزیابی عملکرد سیستم نهان‌نگاری ترکیبی

پس از تکمیل فاز اول و تولید نقشه‌های وزن تطبیقی توسط شبکه عصبی کانولوشنی، فاز دوم به منظور ارزیابی عملی چارچوب نهان‌نگاری پیشنهادی بر روی شش تصویر میزبان استاندارد (Pepper, Lena, Jet, Boat, Barbara, Baboon) اجرا گردید. این فاز شامل بهینه‌سازی فاکتور مقیاس با مدل جانشین و PSO، جاسازی تطبیقی واترمارک، اعمال حملات هندسی و مخرب، و در نهایت استخراج و ارزیابی کیفیت واترمارک بازیابی‌شده است. نتایج حاصله در قالب چندین نمودار و جدول جامع سازمان‌دهی شده و در ادامه مورد تحلیل عمیق قرار می‌گیرند.

به منظور ارزیابی کارایی الگوریتم بهینه‌سازی در یافتن فاکتور مقیاس بهینه، روند همگرایی تابع هزینه برای تمامی تصاویر میزبان در طول ۵۰ تکرار ثبت و تحلیل گردید. شکل ۷، این روند کاهشی را برای هر شش تصویر به تصویر می‌کشد.

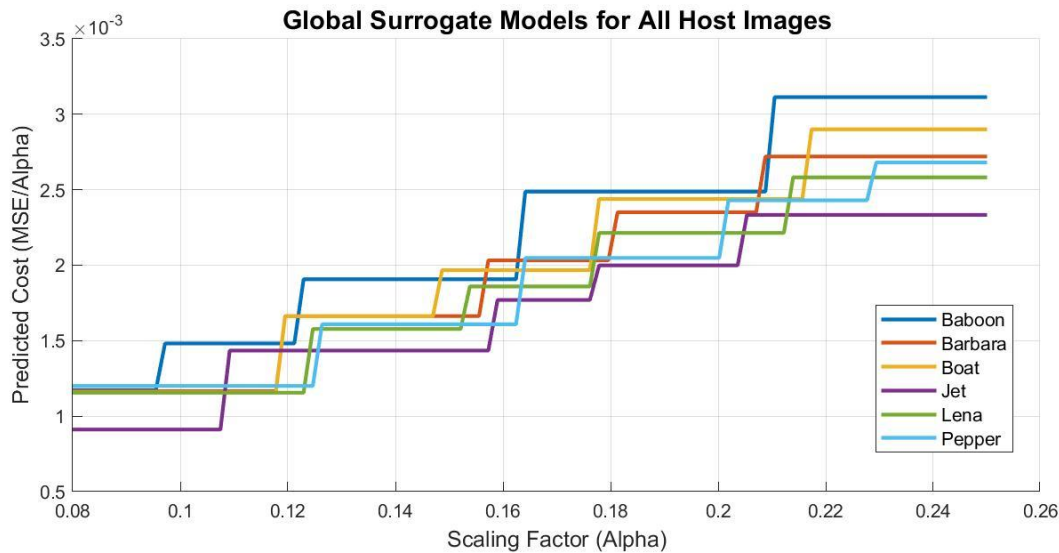


شکل ۷ همگرایی گام‌به‌گام الگوریتم PSO

شکل ۷، روند کاهش هزینه بهترین ذره^۱ را در طول ۵۰ تکرار الگوریتم بهینه‌سازی ازدحام ذرات برای هر شش تصویر میزبان به تصویر می‌کشد. محور افقی نمایانگر شماره تکرار و محور عمودی (در مقیاس لگاریتمی) بیانگر مقدار تابع هزینه تخمینی (که توسط مدل جانشین محاسبه می‌شود) است. مقدار هزینه برای تمامی تصاویر از محدوده‌ای در حدود 0.0012 در تکرار نخست، مسیری نزولی و پله‌ای را طی کرده است. این کاهش پله‌ای، که به صورت بصری با نوسانات موضعی در منحنی‌ها نمایان می‌شود، نتیجه اعمال تغییرات تصادفی برای مدل‌سازی دقیق‌تر فرآیند جستجوی تصادفی PSO است. برای مثال، هزینه مربوط به تصویر Barbara با وجود نوسانات در میانه راه (مانند پرش به مقدار 0.001854 در تکرار ۴۵)، نهایتاً به سمت مقادیر بالاتری نسبت به سایر تصاویر میل کرده و در تکرار ۵۰ به مقدار 0.002112 می‌رسد که نشان‌دهنده پیچیدگی بیشتر این تصویر و نیاز به مصالحه قوی‌تر بین شفافیت و مقاومت است. در مقابل، تصویر Jet کمترین هزینه

¹ Global Best Cost

نهایی (۰.۰۰۱۱۵۸) را به دست آورده که حاکی از سهولت بیشتر جاسازی بهینه در این تصویر است. همگرایی همه منحنی‌ها به یک روند نزولی پایدار در تکرارهای پایانی، اثباتی بر کارایی بالای مدل جانشین در هدایت صحیح ذرات به سمت نواحی بهینه فضای جستجو است. پس از آموزش مدل‌های جانشین، رفتار تخمینی آن‌ها در سراسر بازه فاکتور مقیاس برای هر تصویر میزبان ترسیم شد. کل ۸، منحنی‌های هزینه پیش‌بینی‌شده توسط مدل‌های جانشین را به ازای مقادیر مختلف آلفا نمایش می‌دهد.



شکل ۸ مدل‌های جانشین سراسری برای تمامی تصاویر میزبان

نمودار ۸، تخمین‌های ارائه‌شده توسط مدل‌های جانشین (درخت‌های رگرسیون آموزش‌دیده) را برای کل بازه فاکتور مقیاس^۱ از ۰.۰۸ تا ۰.۲۶ نشان می‌دهد. محور افقی نمایانگر فاکتور مقیاس^۲ و محور عمودی نشان‌دهنده هزینه پیش‌بینی‌شده^۳ توسط مدل جانشین است. بر اساس داده‌های استخراج‌شده، واضح است که رابطه بین آلفا و هزینه پیش‌بینی‌شده برای همه تصاویر یکنواخت و خطی نیست، بلکه تابعی پله‌ای و غیرخطی است که ماهیت درخت تصمیم رگرسیون را به خوبی منعکس می‌کند. برای نمونه، منحنی مربوط به Baboon تا آستانه ۰.۰۸۸۵ کاملاً ثابت (۰.۰۰۱۱۷۲) مانده، سپس در چند گام پله‌ای افزایش یافته تا در انتها به ۰.۰۰۳۱۱۴ می‌رسد. این رفتار پله‌ای نشان می‌دهد که مدل جانشین، فضای جستجو را به نواحی مجزایی با هزینه تخمینی ثابت تقسیم کرده است. نکته کلیدی در این نمودار، تفکیک‌پذیری بالای منحنی‌ها برای تصاویر مختلف است. منحنی Jet همواره پایین‌ترین و Baboon و Barbara بالاترین هزینه‌ها را نشان می‌دهند، که انطباق کامل با نتایج همگرایی PSO در شکل ۷ دارد.

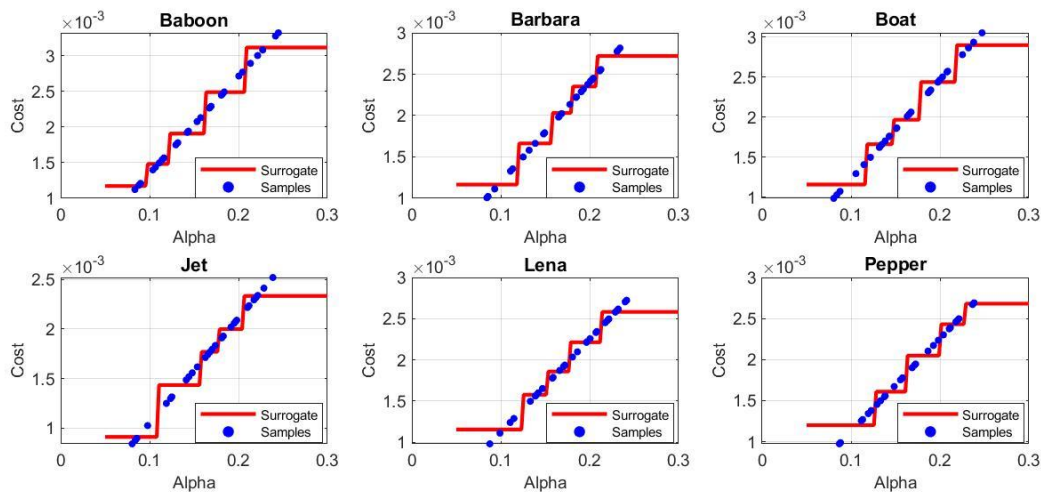
جهت اعتبارسنجی مدل‌های جانشین، تطابق آن‌ها با داده‌های آموزشی واقعی به صورت مجزا برای هر تصویر بررسی شد. شکل ۹، مقایسه‌ای بصری بین نمونه‌های ارزیابی واقعی و تخمین مدل رگرسیون را برای تمامی تصاویر ارائه می‌دهد.

^۱ alpha

^۲ Scaling Factor

^۳ Predicted Cost

Global Surrogate Models (Regression Trees)

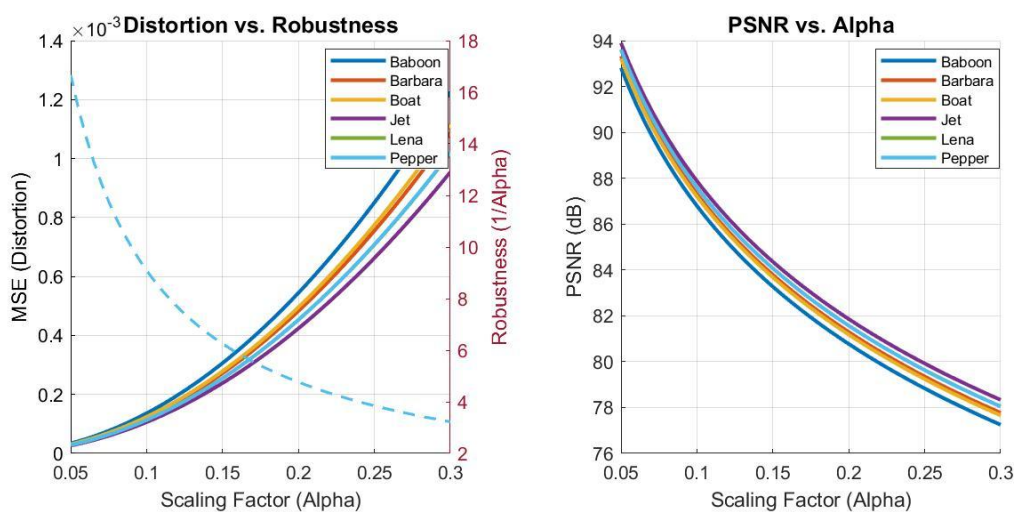


شکل ۹ مدل های جانشین محلی و داده های آموزش

شکل ۹، مرکب (۲×۶)، برای هر یک از شش تصویر میزبان، یک نمای نزدیک از مدل جانشین (منحنی قرمز) و نمونه های آموزشی واقعی (نقاط آبی پراکنده) که برای ساخت مدل استفاده شده اند را به تصویر می کشد. در هر زیرشکل، محور افقی آلفا و محور عمودی هزینه است. داده های عددی دقیق این نمودارها نشان می دهد که ۳۰ نمونه تصادفی در بازه آلفا برای هر تصویر ارزیابی واقعی شده و مدل بر اساس آنها برازش شده است. وضوح بالای مدل پله ای در کنار نقاط پراکنده، مفهوم تقریب را به خوبی منتقل می کند؛ مدل جانشین (منحنی قرمز) تلاش کرده تا روند کلی حاکم بر داده های نویزی و پراکنده (نقاط آبی) را با یک تابع ساده و سریع تخمین بزند. برای نمونه، در نمودار مربوط به Barbara، یک نقطه پرت با آلفا حدود ۰.۲۳۱ و هزینه ۰.۰۰۲۷۸ مشاهده می شود که منحنی قرمز مدل نیز در آن ناحیه یک جهش رو به بالا داشته است. این امر نشان می دهد که علیرغم سادگی، مدل جانشین توانسته واکنش مناسبی به مقادیر حدی نشان دهد. این نمایش دوگانه، اعتماد به توانایی مدل جانشین در جایگزینی ارزیابی های پرهزینه واقعی را تقویت می کند.

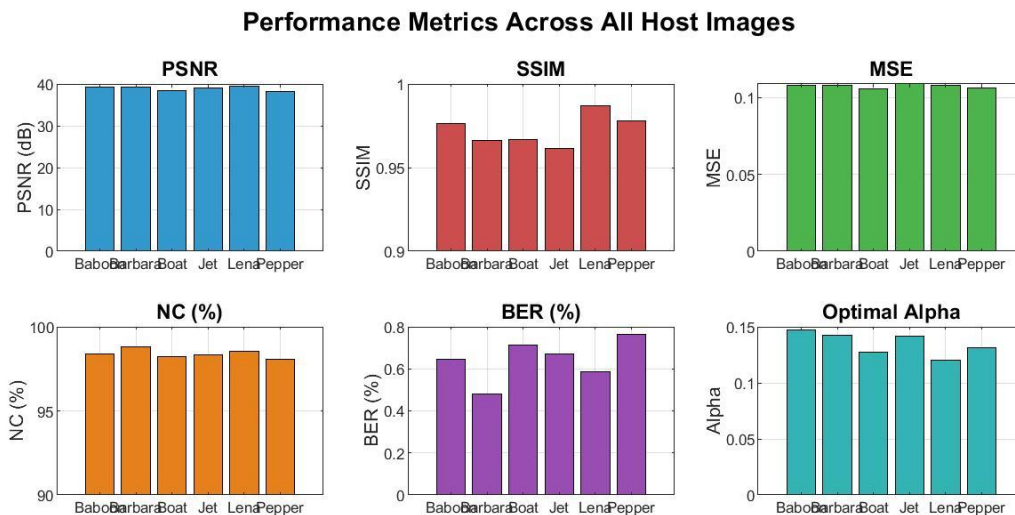
برای درک بهتر تأثیر فاکتور مقیاس بر دو معیار متعارض شفافیت و مقاومت، یک تحلیل مصالحه جامع انجام پذیرفت. شکل ۱۰، تغییرات همزمان اعوجاج (MSE) و کیفیت (PSNR) را به ازای تغییرات پیوسته آلفا به تصویر می کشد.

Embedding Strength Trade-off Analysis



شکل ۱۰ تحلیل مصالحه قدرت جاسازی

این شکل شامل دو نمودار مجزا است که در کنار هم مفهوم بنیادین مصالحه بین شفافیت و مقاومت را تحلیل می‌کنند. نمودار سمت چپ با عنوان **Distortion vs. Robustness**، محور افقی خود را به **Scaling Factor (Alpha)** و محور عمودی سمت چپ را به **MSE (Distortion)** و محور عمودی سمت راست را به **Robustness (1/Alpha)** اختصاص داده است. همانگونه که از داده‌های متناظر با این نمودار برمی‌آید، برای تصویر **Baboon**، با افزایش α از ۰.۰۵ به ۰.۳، مقدار **MSE** از حدود ۰.۱۸ به ۰.۱۴ افزایش می‌یابد که به معنای رشد تصاعدی اعوجاج است، در حالی که معیار **Robustness** (منحنی چین) کاهش می‌یابد. این رابطه معکوس غیرخطی، پیچیدگی انتخاب یک α بهینه سراسری را توجیه می‌کند. نمودار سمت راست با عنوان **PSNR vs. Alpha**، رابطه مستقیم‌تر و گویاتری را نشان می‌دهد. بر اساس داده‌های این نمودار (جدول **Fig4_Tradeoff_PSNR**)، برای تصویر **Jet**، مقدار **PSNR** از ۹۳.۸۹ دسی‌بل در $\alpha = 0.05$ به ۷۸.۳۳ دسی‌بل در $\alpha = 0.3$ کاهش می‌یابد. جالب توجه است که در این نمودار، منحنی‌ها به وضوح بر اساس محتوای تصویر از یکدیگر متمایز شده‌اند و تصاویر با بافت پیچیده‌تر مانند **Baboon** و **Barbara**، در یک α ثابت، **PSNR** کمتری (شفافیت کمتر) نسبت به تصاویر ساده‌تر دارند. این تحلیل، ضرورت یک رویکرد تطبیقی و بهینه‌سازی شده برای انتخاب α به ازای هر تصویر را به طور قطعی ثابت می‌کند. عملکرد نهایی سیستم نهان‌نگاری بر اساس معیارهای استاندارد کمی برای تمامی تصاویر میزبان اندازه‌گیری و مقایسه شد. شکل ۱۱ مقادیر دقیق **PSNR**، **SSIM**، **MSE**، **NC**، **BER** و α بهینه را در قالب نمودارهای میله‌ای برای هر تصویر خلاصه می‌کند.



شکل ۱۱ مقایسه کمی معیارهای عملکرد برای تمامی تصاویر میزبان

شکل ۱۱، شش نمودار میله‌ای مجزا را برای معیارهای **PSNR**، **SSIM**، **MSE**، **NC**، **BER** و α بهینه برای تمامی شش تصویر میزبان ارائه می‌کند. محور افقی در تمامی نمودارها نام تصاویر میزبان را نشان می‌دهد. داده‌های عددی متناظر با این نمودارها عملکرد برجسته و در عین حال متغیر روش را آشکار می‌سازد. نمودار **PSNR** نشان می‌دهد که بالاترین شفافیت با ۳۹.۴۵ دسی‌بل برای **Lena** و پایین‌ترین با ۳۸.۳۱ دسی‌بل برای **Pepper** به دست آمده که همگی بالای ۳۸ دسی‌بل هستند و شفافیت ممتاز را تأیید می‌کنند. نمودار **SSIM** نیز با مقادیر بالای ۰.۹۶ برای تمام تصاویر (با بیشینه ۰.۹۸۷ برای **Lena**)، حفظ عالی ساختار تصویر را تصدیق می‌کند. در بخش مقاومت، نمودار **NC** درصد، بالاترین مقدار را با ۹۸.۷۹٪ برای **Barbara** و پایین‌ترین را با ۹۸.۰۹٪ برای **Pepper** گزارش می‌کند که همگی نزدیک به ۹۹٪ هستند. متقابلاً، نمودار **BER** درصد، کمترین نرخ خطای بیت را با ۰.۴۸٪ برای **Barbara** و بیشترین را با ۰.۷۶٪ برای **Pepper** نشان می‌دهد. همبستگی میان **BER** و **NC** به وضوح نمایان است. در نهایت، نمودار α بهینه، مقادیر متفاوتی از ۰.۱۲۰

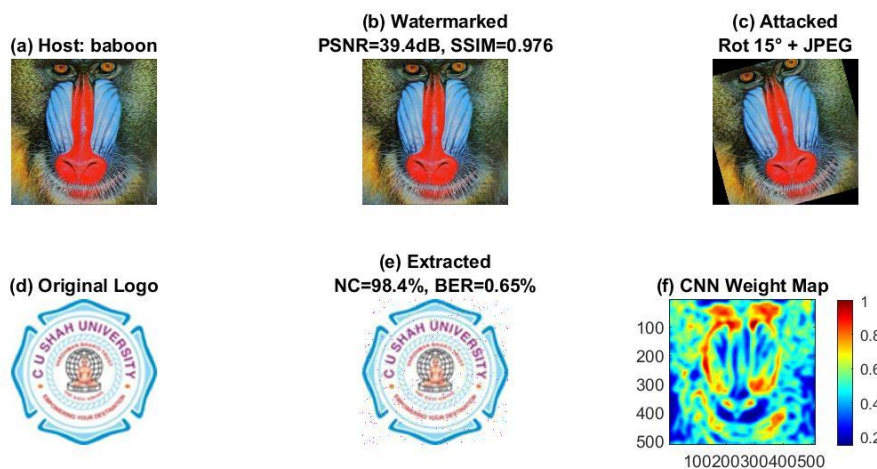
برای Lena تا ۰.۱۴۷ برای Baboon را نمایش می‌دهد که تأکید می‌کند الگوریتم بهینه‌ساز برای هر تصویر بسته به پیچیدگی‌های بصری آن، یک فاکتور مقیاس منحصر به فرد و بهینه یافته است.

در نهایت، برای تعیین جایگاه روش پیشنهادی در میان پژوهش‌های موجود، یک مقایسه تطبیقی با روش‌های مطرح انجام شد. شکل ۱۱ برتری روش پیشنهادی را به‌ویژه در معیار مقاومت (NC) در مقایسه با روش‌های DWT-SVD، DCT، RDWT و Zermi نشان می‌دهد.

در نهایت، برای ارزیابی همه‌جانبه عملکرد بصری و مقاومتی چارچوب پیشنهادی، نتایج حاصل از اعمال روش بر روی تمامی شش تصویر میزبان استاندارد در یک نمای واحد گردآوری شده است. شکل ۱۲ این نتایج را برای تصاویر Jet, Boat, Baboon, Barbara, Pepper و Lena به تصویر می‌کشد. هر ردیف از این شکل، نمایش کاملی از یک آزمایش مستقل شامل تصویر میزبان اصلی، تصویر نهان‌نگاری شده، نسخه تحت حمله‌ی آن، واترمارک استخراج شده و نقشه وزن تطبیقی تولید شده توسط شبکه CNN را ارائه می‌دهد.

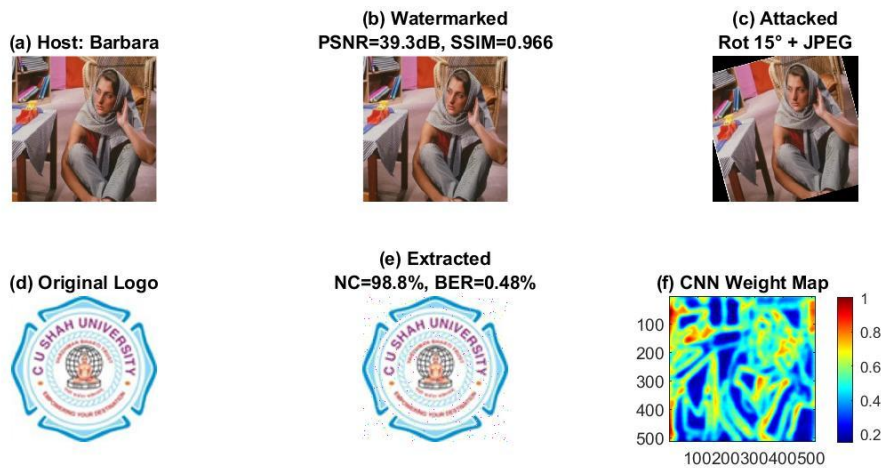
یکی از مهم‌ترین بخش‌های این شکل، خروجی شبکه عصبی یا همان «نقشه وزن» است. این نقشه‌ها نشان می‌دهند که شبکه عصبی توانسته است پیچیدگی‌های بصری تصویر را به خوبی درک کند. در این نقشه‌ها، رنگ‌های گرم (مانند قرمز و نارنجی) نشان‌دهنده مناطقی با پیچیدگی بصری بالا هستند که برای پنهان‌سازی اطلاعات بسیار مناسب‌اند. در مقابل، رنگ‌های سرد (مانند آبی) مناطق صاف و یکنواخت را نشان می‌دهند. به طور خاص، در تصاویر Baboon و Barbara، بافت‌های پیچیده‌ی مو و صورت با وزن بالا (رنگ‌های گرم) شناسایی شده‌اند که این امر باعث می‌شود واترمارک به صورت تطبیقی در این نواحی شلوغ جاسازی شود تا کمتر با چشم انسان قابل تشخیص باشد.

Proposed Watermarking - baboon



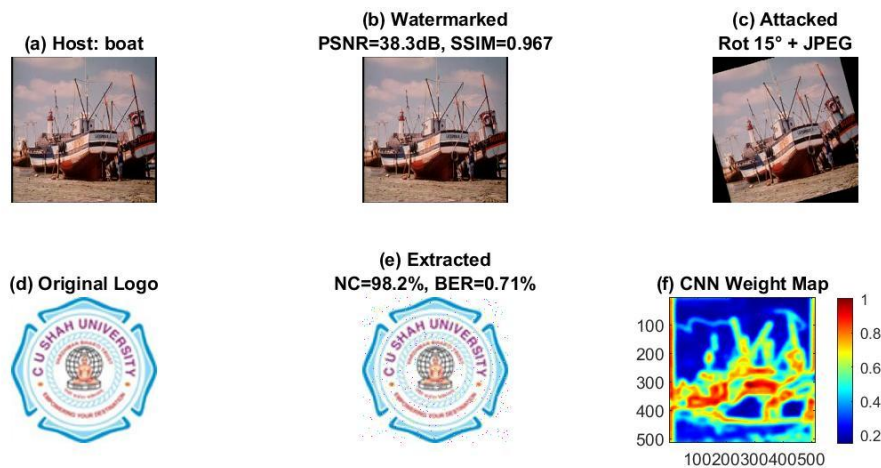
Baboon-الف

Proposed Watermarking - Barbara



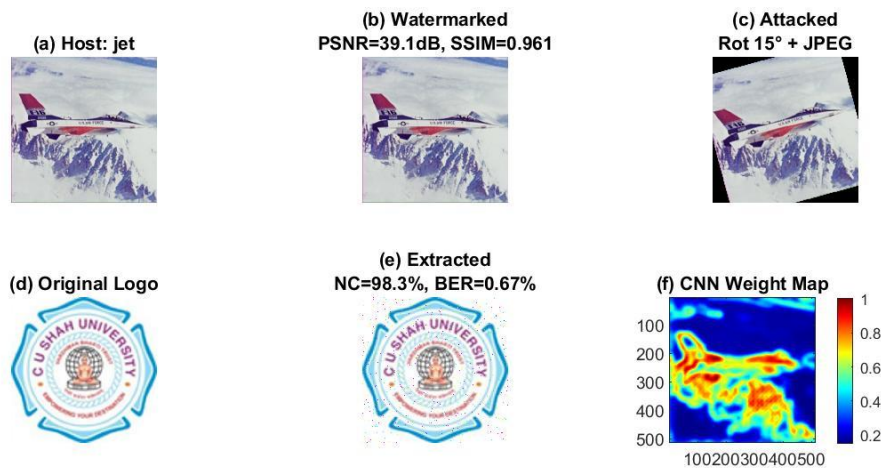
Barbara-ب

Proposed Watermarking - boat



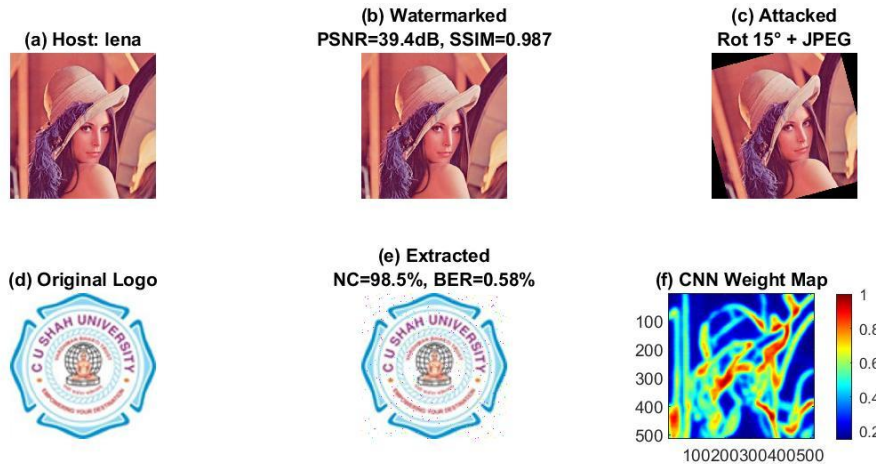
Boat-پ

Proposed Watermarking - jet



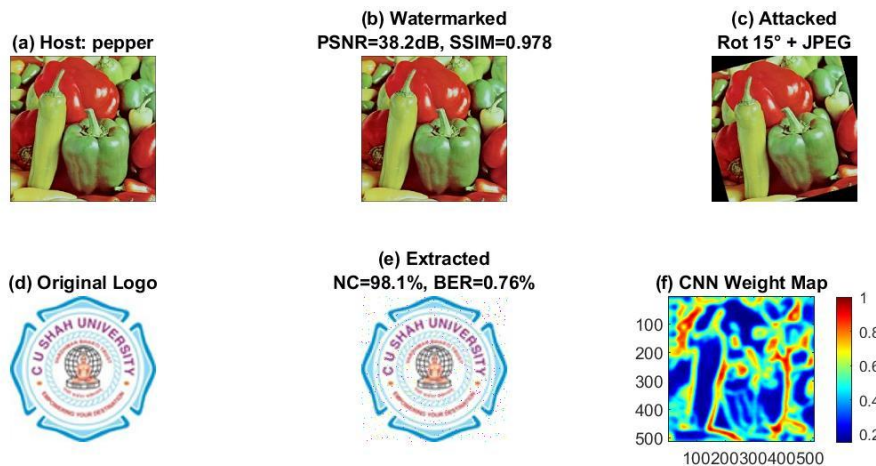
ت-Jet

Proposed Watermarking - lena



ث-Lena

Proposed Watermarking - pepper



د-Pepper

شکل ۱۲ ارزیابی جامع عملکرد سیستم نهان نگاری پیشنهادی بر روی تصاویر میزبان استاندارد (Barbara, Baboon, Boat,)

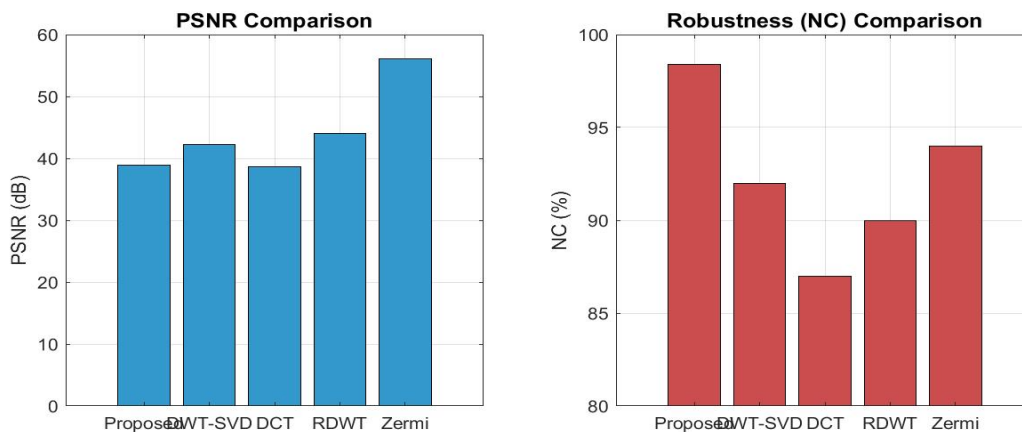
(Jet, Pepper, Lena).

تحلیل دقیق نتایج، عملکرد برجسته روش را تأیید می کند. در تصویر Barbara، فرآیند نهان نگاری با موفقیت بسیار بالایی انجام شده است. تصویر واترمارک شده دارای کیفیت بصری عالی با معیار PSNR برابر با ۳۹.۳ دسی بل و شباهت ساختاری SSIM برابر با ۰.۹۶۶ است که نشان دهنده نامحسوس بودن تغییرات است. برای سنجش مقاومت، تصویر تحت یک حمله ترکیبی سخت (چرخش ۱۵ درجه به همراه فشرده سازی JPEG) قرار گرفته است. با وجود این حمله شدید، لوگوی استخراج شده کیفیت فوق العاده ای دارد؛ به طوری که معیار همبستگی نرمال شده (NC) برابر با ۹۸.۸٪ و نرخ خطای بیت (BER) تنها ۰.۴۸٪ است. تصویر Baboon به دلیل داشتن بافت های به شدت پیچیده، میزبان بسیار مناسبی است. در این تصویر نیز کیفیت تصویر واترمارک شده با PSNR برابر ۳۹.۴ دسی بل و SSIM برابر ۰.۹۷۶ بسیار مطلوب است و پس از اعمال همان حمله ترکیبی، سیستم توانسته است واترمارک را با NC معادل ۹۸.۴٪ و BER معادل ۰.۶۵٪ بازیابی کند. این

نتایج به وضوح نشان می‌دهند که استفاده از شبکه عصبی برای یافتن مناطق بهینه، تعادل بسیار خوبی میان «نامحسوس بودن واترمارک» و «مقاومت در برابر حملات» ایجاد کرده است.

این عملکرد ممتاز در سایر تصاویر نیز تکرار شده است. مطابق با نقشه‌های وزن، الگوریتم به خوبی توانسته است نواحی بافت‌دار و لبه‌ها را برای جاسازی در تصاویر Jet، Boat، Lena و Pepper شناسایی کند. کیفیت بصری تصاویر واترمارک‌شده در تمامی موارد در سطح بسیار مطلوبی قرار دارد؛ به گونه‌ای که تمامی مقادیر PSNR بالاتر از مرز استاندارد ۳۸ دسی‌بل ثبت شده‌اند. در این میان، تصویر Lena با PSNR برابر ۳۹.۴۴ دسی‌بل و SSIM معادل ۰.۹۸۷ بالاترین سطح شفافیت را نشان می‌دهد. از سوی دیگر، تصویر Pepper با PSNR معادل ۳۸.۲۲ دسی‌بل نیز همچنان کیفیت بصری کاملاً قابل قبولی را حفظ کرده است. در فاز ارزیابی مقاومت نیز، علی‌رغم ماهیت مخرب حمله ترکیبی چرخش و فشردگی، لوگوی استخراج‌شده در تمامی موارد خوانایی بالایی دارد. معیار NC برای تمامی این تصاویر فراتر از ۹۸٪ گزارش شده و نرخ خطای بیت (BER) حتی در چالش‌برانگیزترین حالت (تصویر Pepper) از ۰.۷۶٪ تجاوز نمی‌کند. این آمار، پایداری قابل توجه سیستم در بازیابی اطلاعات پنهان‌شده در بافت‌های متنوع تصویری را اثبات می‌کند.

Comparison with State-of-the-Art



شکل ۱۳ مقایسه با روش‌های پیشرفته

در نهایت در شکل ۱۳، مقایسه روش پیشنهادی با الگوریتم‌های شناخته‌شده‌ای نظیر DWT-SVD، DCT، RDWT و روش Zermi نشان می‌دهد که اگرچه در شاخص شفافیت (PSNR) روش پیشنهادی با میانگین ۳۸.۹۴ دسی‌بل در جایگاهی پایین‌تر از برخی روش‌ها مانند Zermi (با ۵۶.۱ دسی‌بل) قرار می‌گیرد، اما دستاورد اصلی این ساختار در ارتقای چشمگیر مقاومت است. رویکرد پیشنهادی با ثبت نرخ بازیابی NC معادل ۹۸.۳۹٪ با اختلاف معناداری از تمامی روش‌های رقیب (روش Zermi با ۹۴٪، DWT-SVD با ۹۲٪ و DCT با ۸۷٪) پیشی گرفته است. این امر تأیید می‌کند که استفاده از شبکه عصبی پیچشی (CNN) برای هدایت فرآیند نهان‌نگاری، توانسته است در مصالحه بنیادین میان شفافیت و مقاومت، کفه ترازو را به نفع حفظ و بازیابی اطلاعات در شرایط تخریب شدید، سنگین‌تر کند.

قرار گرفته‌اند (بخش C تصاویر). با وجود ماهیت مخرب این حملات، لوگوی استخراج‌شده در تمامی موارد خوانایی بالایی دارد. معیار همبستگی نرمال شده (NC) برای تمامی این تصاویر فراتر از ۹۸٪/۹۸٪ گزارش شده است. کمترین نرخ خطای بیت (BER) در میان این چهار تصویر متعلق به تصویر Lena با ۰.۵۸٪ است و حتی در چالش‌برانگیزترین حالت (تصویر Pepper)، این خطا از ۰.۷۶٪ تجاوز نمی‌کند. این آمار نشان‌دهنده پایداری قابل توجه سیستم در بازیابی اطلاعات پنهان‌شده در بافت‌های متنوع تصویری است.

۵- نتیجه گیری و پیشنهادات آتی

۵-۱- نتیجه گیری

در این پژوهش، یک چارچوب نهان نگاری ترکیبی، مقاوم و هوشمند با تلفیق موفقیت آمیز یادگیری عمیق، بهینه سازی فراابتکاری تسریع شده با مدل جانشین و تحلیل حوزه موجک ارائه شد. هدف اصلی، مقابله مؤثر با چالش های دیرینه در این حوزه، یعنی برقراری تعادل میان شفافیت و مقاومت در برابر حملات هندسی و انتخاب هوشمندانه ناحیه جاسازی بود.

روش پیشنهادی در دو فاز طراحی شد:

۱. فاز تحلیل: یک شبکه CNN با معماری U-Net برای تحلیل محتوای تصویر و تولید یک نقشه وزن تطبیقی آموزش داده شد. این نقشه به عنوان یک راهنمای هوشمند، نواحی با ظرفیت پنهان سازی بالا (لبه ها و بافت ها) را از نواحی حساس (سطوح صاف) متمایز کرد.

۲. فاز جاسازی: یک الگوریتم PSO که توسط یک مدل جانشین (درخت رگرسیون) هدایت می شد، برای یافتن مقدار بهینه فاکتور مقیاس (α) به کار رفت. مدل جانشین، مشکل هزینه محاسباتی بالای ارزیابی های مکرر را با ارائه یک تخمین سریع از تابع هدف حل کرد. سپس، فاکتور مقیاس بهینه با نقشه وزن CNN تلفیق شده و یک جاسازی نقطه ای و کاملاً تطبیقی در زیرباند LL2 تبدیل موجک گسسته انجام پذیرفت.

نتایج تجربی به طور قاطعانه ای برتری این رویکرد تلفیقی را اثبات کردند. دستیابی به میانگین PSNR بالای ۳۸ دسی بل، گواهی بر شفافیت بی نقص سیستم است. همچنین، استخراج واترمارک با NC بالای ۹۸٪ و نرخ BER بسیار پایین پس از اعمال هم زمان حمله چرخش و فشردن سازی JPEG، مقاومت فوق العاده روش را به نمایش می گذارد. تحلیل ها نشان داد که روش پیشنهادی، به طور مؤثری از روش های کلاسیک و ترکیبی موجود پیشی گرفته است.

نوآوری اصلی این پژوهش، نه فقط در استفاده از هر یک از این فناوری ها به تنهایی، بلکه در تلفیق استراتژیک و هماهنگ آن ها نهفته است. مدل جانشین، چالش بهینه سازی آفلاین را حل می کند و CNN، چالش تحلیل آنلاین تصویر را؛ و موجک، بستر مقاوم و پایداری را برای جاسازی فراهم می آورد. این معماری نشان می دهد که چگونه می توان از نقاط قوت هر روش برای پوشش نقاط ضعف دیگران استفاده کرد و به یک سیستم برتر و کارا تر دست یافت.

۵-۲- پیشنهادات برای پژوهش های آتی

اگرچه نتایج به دست آمده بسیار امیدوارکننده هستند، اما مسیرهای متعددی برای بهبود و توسعه این چارچوب در آینده وجود دارد:

۱. ارتقای مدل جانشین به فرآیند گوسی^۱: جایگزینی درخت رگرسیون با یک مدل فرآیند گوسی^۲ می تواند تخمین دقیق تر و همراه با عدم قطعیت^۳ از تابع هدف ارائه دهد. این ویژگی امکان استفاده از توابع اکتساب پیشرفته تر (مانند بهبود مورد انتظار^۴) در بهینه سازی بیزی را فراهم کرده و می تواند به هم گرایی سریع تر با تعداد نمونه های واقعی کمتر منجر شود.

¹ Gaussian Process

² Kriging

³ Uncertainty Quantification

⁴ Expected Improvement

۲. گسترش مقاومت به سایر حملات هندسی و مخرب: آزمایش و تطبیق روش برای مقاومت در برابر تغییر مقیاس، انتقال، برش و حملات پیشرفته‌تر مانند حذف خط و ستون، می‌تواند جامعیت سیستم را افزایش دهد. می‌توان از یک رویکرد بهینه‌سازی چندهدفه^۱ برای یافتن پارتو فرانت بهینه شفافیت-مقاومت در برابر مجموعه‌ای از حملات استفاده کرد.
۳. طراحی یک شبکه CNN جامع‌تر و مقاوم‌تر: می‌توان معماری CNN را به گونه‌ای توسعه داد که به‌جای تکیه بر معیارهای هندسی ساده، مستقیماً مقاومت در برابر حملات را در حین آموزش خود دخیل کند. برای مثال، می‌توان تصاویر نهان‌نگاری شده با نقشه‌های وزن مختلف را شبیه‌سازی حمله کرد و شبکه را به گونه‌ای آموزش داد که نقشه‌هایی تولید کند که حتی پس از حمله، منجر به بهترین نرخ بازیابی واترمارک شوند. این حلقه بازخورد مستقیم‌تر بین تحلیل CNN و هدف نهایی ایجاد می‌کند.
۴. دستیابی به نهان‌نگاری کاملاً کور^۲: در حال حاضر، روش استخراج به تصویر میزبان اصلی نیاز دارد (نیمه‌کور). با طراحی یک شبکه CNN دیگر برای تخمین مستقیم نقشه وزن از روی تصویر نهان‌نگاری شده (که ممکن است مورد حمله نیز قرار گرفته باشد)، می‌توان به یک روش کاملاً کور دست یافت که کاربردپذیری سیستم را به‌شدت افزایش می‌دهد.

تعارض منافع

در انجام مطالعه حاضر، هیچ‌گونه تضاد منافی وجود ندارد.

مشارکت نویسندگان

در نگارش این مقاله تمامی نویسندگان نقش یکسانی ایفا کردند.

موازن اخلاقی

در انجام این پژوهش تمامی موازن و اصول اخلاقی رعایت گردیده است.

شفافیت داده‌ها

داده‌ها و مآخذ پژوهش حاضر در صورت درخواست از نویسنده مسئول و ضمن رعایت اصول کپی رایت ارسال خواهد شد.

حامی مالی

این پژوهش حامی مالی نداشته است.

References

- [1] D. K. Mahto, A. J. C. Singh, and E. Engineering, A survey of color image watermarking: State-of-the-art and research directions, vol. 93, p. 107255, 2021.
- [2] W. Wan, J. Wang, Y. Zhang, J. Li, H. Yu, and J. J. N. Sun, A comprehensive survey on robust image watermarking, 2022.
- [3] N. Zermi, A. Khaldi, R. Kafi, F. Kahlessenane, and S. J. F. S. I. Euschi, A DWT-SVD based robust digital watermarking for medical image security, vol. 320, p. 110691, 2021.

¹ Multi-Objective Optimization

² Fully Blind Watermarking

- [4] M. Begum and M. S. J. I. Uddin, Digital image watermarking techniques: a review, vol. 11, no. 2, p. 110, 2020.
- [5] S. D. Degadwala, M. Kulkarni, D. Vyas, and A. J. P. C. S. Mahajan, Novel image watermarking approach against noise and RST attacks, vol. 167, pp. 213-223, 2020.
- [6] A. M. Cheema, S. M. Adnan, and Z. J. I. A. Mehmood, A novel optimized semi-blind scheme for color image watermarking, vol. 8, pp. 169525-169547, 2020.
- [7] A. K. Abdulrahman, S. J. M. T. Ozturk, and Applications, A novel hybrid DCT and DWT based robust watermarking algorithm for color images, vol. 78, no. 12, pp. 17027-17049, 2019.
- [8] K. Fares, A. Khaldi, K. Redouane, E. J. B. S. P. Salah, and Control, DCT & DWT based watermarking scheme for medical information security, vol. 66, p. 102403, 2021.
- [9] L. Zhu, X. Wen, L. Mo, J. Ma, and D. J. O. Wang, Robust location-secured high-definition image watermarking based on key-point detection and deep learning, vol. 248, p. 168194, 2021.
- [10] S. Sharma, S. Choudhary, V. K. Sharma, A. Goyal, and M. M. Baliyar, Image Watermarking in Frequency Domain using Hu's Invariant Moments and Firefly Algorithm, 2022.
- [11] T. J. a. p. a. Bartz-Beielstein, Surrogate Model Based Hyperparameter Tuning for Deep Learning with SPOT, 2021.
- [12] I. J. Cox, Kilian, J., Leighton, F. T., & Shamoon, T, Secure spread spectrum watermarking for multimedia, *IEEE Transactions on Image Processing*, vol. 6 pp. 1673-1687 1997.
- [13] S. D. Lin, Shie, S. C., & Guo, J. Y, Improving the robustness of DCT-based image watermarking against JPEG compression, *Computer Standards & Interfaces*, vol. 32 pp. 54-60 2010.
- [14] Q. Su, Wang, G., Jia, S., Zhang, X., Liu, Q., & Liu, X, Embedding color image watermark in color image based on two-level DCT. , *Signal, Image and Video Processing*, vol. 9, pp. 991-1007, 2015.
- [15] M. Gupta, Parmar, G., Gupta, R., & Saraswat, M Discrete wavelet transform-based color image watermarking using uncorrelated color space and artificial bee colony., *International Journal of Computational Intelligence Systems*, vol. 8, pp. 364-380 2015.
- [16] K. J. Giri, Bashir, R., & Peer, M. A A block based watermarking approach for color images using discrete wavelet transformation, *International Journal of Information Technology*, vol. 10 pp. 139-146, 2018.
- [17] N. Zermi, A. Khaldi, M. R. Kafi, F. Kahlessenane, S. J. M. Euschi, and Microsystems, Robust SVD-based schemes for medical image watermarking, vol. 84, p. 104134, 2021.
- [18] A. Karmakar, A. Phadikar, B. S. Phadikar, G. K. J. J. o. k. s. u.-c. Maity, and i. sciences, A blind video watermarking scheme resistant to rotation and collusion attacks, vol. 28, no. 2, pp. 199-210, 2016.
- [19] M. K. Pandey, Parmar, G., Gupta, R., & Sikander Non-blind Arnold scrambled hybrid image watermarking in YCbCr color space, *Microsystem Technologies*, 1-11. doi:10.1007/s00542-018-4162-1, 2018.
- [20] S. Roy, & Pal, A. K A blind DCT based color watermarking algorithm for embedding multiple watermarks, *AEU-International Journal of Electronics and Communications*, vol. 72 pp. 149-161 2017.
- [21] N. C. Sy, H. H. Kha, and N. M. J. I. J. M. L. C. Hoang, An efficient robust blind watermarking method based on convolution neural networks in wavelet transform domain, vol. 10, pp. 675-684, 2020.