

Imágenes tomadas de Canva educativo

Algoritmo de marcado de agua robusto: autenticación de imágenes mediante redes neuronales siamesas y características híbridas espacio-frecuencia

Robust zero-watermarking algorithm for image authentication using hybrid spatial-frequency features and siamese neural networks

Rodrigo Eduardo Arevalo-Ancona, Eduardo Frago-so-Navarro, Manuel Cedillo-Hernández*

RESUMEN

El uso y distribución de archivos digitales, impulsado por los avances en las tecnologías de la información, ha generado la necesidad de desarrollar sistemas para la protección de los derechos de autor. En el contexto de las imágenes digitales, es fundamental minimizar los riesgos de seguridad asociados con la distribución no autorizada y garantizar la integridad de la información visual. El objetivo de este trabajo fue desarrollar un algoritmo de marcado de agua libre de distorsiones, diseñado para la autenticación de la propiedad del usuario y la recuperación de imágenes originales en escala de grises en caso de su manipulación. El método propuesto utilizó una red neuronal siamesa con una arquitectura de dos ramas: una que aprende características frecuenciales a partir de los coeficientes de la transformada de wavelet discreta y otra que extrae características espaciales de la otra imagen. Adicionalmente, se integró una red neuronal entrenada con características espaciales para reconstruir una versión en escala de grises de la imagen original, natural a color, tras una manipulación. El método propuesto demostró su eficacia en los tiempos de procesamiento, precisión y proceso de recuperación de la marca de agua para la verificación de la propiedad del usuario y la autenticación de imágenes manipuladas. Destacando sus mejoras de robustez, frente a distorsiones geométricas como rotación, traslación, transformación, afín, recorte y escalado, así como la combinación de algunas distorsiones. El autoencoder entrenado conserva una alta fidelidad en la reconstrucción de imágenes en escala de grises, alteradas por manipulaciones combinadas con otras distorsiones, por lo que demostró ser una solución eficaz para la autenticación y protección de los derechos de autor en imágenes digitales.

PALABRAS CLAVE: marcado de agua libre de distorsiones, aprendizaje profundo, red neuronal siamesa, autenticación de propietario, seguridad en imágenes.

ABSTRACT

The use and distribution of digital files has increased due to advancements in information technologies. This has created the need for the development of copyright protection systems. In the context of digital images, it is essential to minimize security risks associated with unauthorized distribution and ensure the integrity of visual information. The objective of this paper was to develop a robust zero-watermarking algorithm designed for user ownership verification and the recovery of original images in grayscale in cases of manipulation. The proposed method employed a Siamese neural network with a architecture model consisted of two branches: one learns frequency-domain features from discrete wavelet transform coefficients, and other that extract spatial features from the other image. Additionally, a neural network trained with spatial features was used to reconstruct a grayscale version of the original image after tampering. The proposed method proved its using natural color images demonstrated the high effectiveness in processing times, precision and of the watermark recovery process for verifying user ownership and accurately authenticating manipulated images. Highlighted its improvements robustness when encountering geometrical distortions such as rotation, translation, transformation, affine, clipping and scaling, as well as the combination of some distortions. The trained autoencoder preserves high fidelity in grayscale image reconstruction that has suffered alterations by combining manipulations with other distortions. The proposed algorithm proved to be an effective solution for the authentication and protection of copyright in digital images.

KEYWORDS: zero-watermarking, deep learning, siamese neural network, ownership authentication, image security.

*Correspondencia: mcedilloh@ipn.mx/ Fecha de recepción: 17 de enero de 2025/ Fecha de aceptación: 27 de junio de 2025/ Fecha de publicación: 18 de agosto de 2025.

Instituto Politécnico Nacional, Escuela Superior de Ingeniería Mecánica y Eléctrica, Unidad Culhuacán, Sección de Estudios de Posgrado e Investigación, avenida Sta. Ana, núm. 1000, San Francisco Culhuacán, Culhuacán CTM V, Coyoacán, Ciudad de México, México, C. P. 04440.

INTRODUCCIÓN

Los avances recientes en las tecnologías de la información han provocado un notable aumento en la creación y distribución de archivos digitales, especialmente imágenes digitales (Ikbal y Gopikakamari, 2022). Esta creciente demanda ha generado la necesidad de desarrollar soluciones eficaces para proteger, preservar y verificar la autenticidad de las imágenes (Shamia y col., 2023). Por esta razón, se han desarrollado técnicas de marcado de agua para proteger las imágenes digitales y evitar falsificaciones o usos no autorizados (Ami ni y col., 2018; Wu y col., 2020).

Las marcas de agua utilizadas contienen información derivada de la imagen o del propietario, permitiendo identificar su propiedad (Lee y col., 2019; Liu y col., 2020; Shang y col., 2023). Las metodologías convencionales insertan una señal en la imagen para proteger su propiedad intelectual, pero generan distorsiones que afectan su calidad visual (Gong y col., 2022). Para solucionar esta limitación, se han desarrollado algoritmos de marcado de agua libre de distorsiones.

Los métodos libres de distorsiones emplean un componente adicional, el código de usuario, que permite autenticar la imagen sin modificarla ni comprometer su calidad visual. Este código se almacena y se genera asociando características de la imagen con una marca de agua binaria mediante una operación lógica (Liu y Zhang, 2021). Durante la autenticación, el código de usuario se combina con las características de la imagen, para recuperar la marca de agua, garantizando la autenticación de la propiedad (Dong y col., 2023b; Khafaga y col., 2023).

Zhang (2021) presentó un sistema libre de distorsiones utilizando el coeficiente DC de la transformada discreta del coseno (DCT, por sus siglas en inglés: discrete cosine transform) para generar una matriz de características. El uso de características frecuenciales (asociadas con la distribución de la intensidad de píxeles) mejora la robustez del sistema ante dis-

torsiones geométricas. Posteriormente, utiliza el algoritmo de descomposición en valores singulares (SVD, por sus siglas en inglés: singular vector decomposition), para extraer las características de la imagen. Este método garantiza la extracción de características robustas. Li y col. (2023b) propusieron un algoritmo para imágenes médicas, combinando los algoritmos del descriptor de orientación y rotación (ORB, por sus siglas en inglés: oriented fast and rotated brief) y DCT para crear una matriz de características y fusionarla con una marca de agua encriptada. Esta técnica identifica características locales y globales de la imagen.

La integración de modelos de aprendizaje profundo en los sistemas de marcado de agua mejora la robustez al identificar características únicas de la imagen, incluso cuando la imagen tenga distorsiones como transformaciones geométricas o procesamiento (Zhong y Shih, 2019; Jing, 2020; Solorzano y Tsai, 2023).

Las redes neuronales aprenden patrones complejos de las imágenes, lo que permite autenticarlas incluso después de haber sufrido diversas modificaciones. Han y col. (2021) utilizaron un modelo preentrenado de VGG19 para extraer los mapas de características, en el que, posteriormente, aplicaron la transformada discreta de Fourier (DFT, por sus siglas en inglés: discrete Fourier transform). El uso de DFT mejoró la robustez frente a manipulaciones geométricas. Darwish y col. (2023) desarrollaron un método usando las características extraídas de la imagen mediante la red neuronal VGG19. Huang y col. (2023) emplearon el modelo VGG para construir una matriz de características, utilizada para la protección de la imagen, resaltando las características detectadas. Gong y col. (2022) obtuvieron los mapas de características del modelo residual DenseNet, los cuales combinaron con los coeficientes de la DCT. El uso de un modelo DenseNet ayudó a identificar características de datos no vistos. Dong y col. (2023a) propusieron un sistema basado en NasNet Mobile, al que añadieron una capa de regresión en lugar de la capa de clasificación para

obtener 128 características, a las cuales se aplicó DCT. Xiang y col. (2023) emplearon un modelo neuronal ResNet como extractor de características. La red neuronal se entrena con características de diferentes filtros aplicados a la imagen. A pesar de la efectividad de los métodos de aprendizaje profundo para extraer características robustas, la mayoría de los métodos se enfocan en distorsiones geométricas o de procesamiento de señales.

El objetivo del presente trabajo fue desarrollar un sistema para la recuperación de imágenes en escala de grises, con el fin de determinar si la imagen ha sido distorsionada y autenticar al propietario, basado en un algoritmo de marcado de agua libre de distorsiones que utiliza una red neuronal siamesa como extractor de características.

MATERIALES Y MÉTODOS

Se utilizó el conjunto de datos pertenecientes a MICC-F220 (Amerini y col., 2011) para entrenar y validar la autenticación de usuario; y se empleó el conjunto COCO 2017 por su diversidad de imágenes para evaluar la reconstrucción de la imagen (Lin y col., 2015). Se utilizó un código QR de 512×512 píxeles como marca de agua. El algoritmo se implementó en una computadora Alienware M16 R2 (marca registrada de Dell, Round Rock, Texas, Estados Unidos) con un sistema con GPU NVIDIA GTX 4050 y un procesador Intel Core Ultra 7 155H (1.4 GHz) con Windows 11, usando Pytorch versión 2.4.1. El equipo fue fabricado en China o México, dependiendo de la planta de producción.

El método propuesto se dividió en cinco etapas:

1. Preprocesamiento de la imagen. Se aplicó aumento de datos a la rama espacial y transformada discreta de wavelet (DWT, por sus siglas en inglés: discrete wavelet transform) para obtener los coeficientes de frecuencia baja (LL, por sus siglas en inglés: low-low). Se redimensionó la imagen a una dimensión de 128×128 píxeles para asegurar compatibilidad con la red neuronal siamesa de dos ra-

mas, que procesa características espaciales y frecuenciales.

2. Entrenamiento de la red neuronal. La red aprendió patrones invariantes en las características espaciales y frecuenciales. La combinación de ambas ramas creó una matriz de características.

3. Generación del código de usuario. Se creó la matriz de características y se vinculó lógicamente con la marca de agua, para verificar al usuario de manera eficaz.

4. Recuperación de la marca de agua. La red neuronal extrajo la matriz de características de la imagen distorsionada, y al combinarla con el código de usuario, se recuperó la marca de agua, incluso en imágenes manipuladas.

5. Reconstrucción de imagen sin distorsiones. Se utilizó una red neuronal artificial que aprendió a comprimir y reconstruir datos de entrada denominada autoencoder, para reconstruir la imagen distorsionada en escala de grises, verificando su autenticidad.

Preprocesamiento de las imágenes

En esta etapa, se aplicaron dos procesamiento de imágenes para generar las entradas para cada rama de la red neuronal. Primero, se realizó un proceso de aumento de datos para introducir distorsiones geométricas y de procesamiento de imágenes a las imágenes originales O_k del conjunto de datos, donde $k = 0, \dots, L$, generando versiones modificadas A_k , permitiendo detectar características invariantes, mejorando su capacidad para aprender representaciones robustas.

Las técnicas de aumento de datos empleadas incluyen: rotaciones de imágenes de 0° a 360° , con incrementos de 15° . Traslaciones de imágenes en el eje x, eje y, o ambos, con desplazamientos que van de 10 a 150 píxeles en pasos de 20 píxeles. Escalado de imágenes a dimensiones de 64×64 , 128×128 , 512×512 y 1024×1024 píxeles. Recorte de imágenes desde las esquinas y el centro,

con tamaños de 50 x 50, 100 x 100 y 150 x 150 píxeles. Aplicación de diferentes filtros, incluyendo filtros promedio, filtros gaussianos, filtros medianos y desenfoque, con tamaños de núcleo de 5 x 5, 7 x 7 y 11 x 11 píxeles. Introducción de ruido sal y pimienta con densidades del 5 %, 9 % y 15 %. Compresión JPEG con factores de calidad de imagen de 30, 50, 70 y 90 píxeles.

Posteriormente, se aplicó el primer nivel de DWT para extraer los coeficientes de frecuencia LL, denotados como F_k . Estos coeficientes se seleccionan porque contienen información de baja frecuencia, y permiten eliminar características redundantes, así como ruido que contenga la imagen, lo cual aumenta la robustez frente a distorsiones, como compresión, transformaciones geométricas y adición de ruido. Las imágenes F_k y A_k se redimensionaron a un tamaño estandarizado de 128 x 128 píxeles. Este redimensionamiento garantiza dimensiones de entrada consistentes, simplificando el entrenamiento de la red neuronal y la estructura de los datos.

La incorporación de los coeficientes LL de la DWT en el entrenamiento de la red neuronal garantizó una recuperación precisa de la marca de agua para la autenticación de propiedad, incluso cuando la imagen ha sido distorsionada. Al integrar estos coeficientes, se estableció una relación estable entre las características espaciales y frecuenciales de la imagen. Dicha relación permite a la red neuronal identificar patrones invariantes, para la creación de una matriz de características que se mantiene similar tanto cuando la imagen está intacta como cuando está distorsionada, lo que asegura la precisión en la recuperación de la marca de agua y la verificación de la propiedad. La siguiente fase consistió en entrenar la red neuronal utilizando tanto las imágenes A_k como los coeficientes F_k .

Entrenamiento de la red neuronal

En esta etapa, la red neuronal aprende patrones invariantes a partir de los datos. En el contexto del marcado de agua libre de distorsiones,

las redes neuronales siamesas son utilizadas como un extractor de características, que son cruciales para crear el código de usuario y autenticar la imagen (Arevalo-Ancona y col., 2024).

La red neuronal siamesa AlexNet consta de dos ramas idénticas con los mismos parámetros, y está diseñada para aprender la correlación de patrones entre la imagen original y la imagen distorsionada, lo que permite identificar patrones comunes entre las entradas (Vizváry y col., 2019; Wiggers y col., 2019; Aydemir y col., 2022; Li y col., 2023a).

El extractor de características basado en la red neuronal siamesa se entrena utilizando una función de pérdida contrastiva $CLoss$. Esta función de pérdida permite a la red neuronal aprender representaciones invariantes de las imágenes de entrada, que son esenciales para construir matrices de características similares tanto para las imágenes originales como para las distorsionadas. La pérdida contrastiva minimiza la distancia entre los vectores de salida, asegurando que sus representaciones permanezcan similares a pesar de las distorsiones. Al mismo tiempo, maximiza la distancia entre pares de imágenes disímiles, para distinguir entre características no relacionadas (Ye y col., 2025). Esta función de pérdida permite a la red neuronal aprender una representación semántica de las imágenes al capturar efectivamente patrones similares de los coeficientes derivados mediante la DWT y las características espaciales.

La $CLoss$ se define como (1):

$$CLoss = \frac{1}{N} \sum_{k=1}^N [(1 - y_k)(d) + y_k(\max(0, m - d))] \quad (1)$$

Donde:

N = el lote de imágenes

y = valor de la característica analizada

$y_k = 0$ cuando A_k (imagen) está distorsionada

$y_k = 1$ cuando A_k (imagen) no ha sido modificada

El término m ($m = 2$) define la distancia mínima d (2) entre los vectores de salida: V_{out1} y V_{out2} .

$$d = \sqrt{V_{out1} - V_{out2}} \quad (2)$$

Donde:

V_{out1} y V_{out2} = vectores de salidas de la red neuronal siamesa

Para reforzar la autenticación de imágenes distorsionadas, se utilizó un autoencoder para reconstruir la imagen sin distorsiones en escala de grises. Su codificador comparte la arquitectura de la rama espacial de la red siamesa, proporcionando características invariantes y consistentes. El decodificador, mediante convoluciones transpuestas, reconstruyó una imagen similar a la original. Este diseño optimizó la eficiencia computacional mejorando la reconstrucción de las imágenes distorsionadas.

Las imágenes A_k se utilizaron como entrada para el autoencoder, donde el codificador comprimió estas imágenes en una representación latente y el decodificador generó la imagen reconstruida. Para lograr un aprendizaje robusto del espacio latente, se incorporó una reparametrización (3):

$$r_z = \mu + \epsilon(\sigma) \quad (3)$$

Donde:

r_z = espacio latente

ϵ = salida

μ = media

σ = desviación estándar

La reparametrización en el autoencoder permitió introducir variabilidad controlada en el espacio latente, facilitando un aprendizaje más robusto y estable. El codificador generó la salida ϵ y se obtuvo su media (μ) y desviación estándar (σ) generando un espacio latente (r_z), modelado como una distribución de probabilidad continua. Esto aseguró que el modelo pudiera capturar mejor las variaciones de las imágenes distorsionadas y mejorar la calidad de la reconstrucción, generando imágenes más fieles a las originales.

La función de pérdida V_{Loss} se obtiene mediante la siguiente fórmula (4):

$$V_{Loss} = L_{recon} + \beta L_{KL} \quad (4)$$

Donde:

L_{recon} = pérdida de reconstrucción

β = Parámetro equilibra la calidad en reconstrucción de la imagen y el espacio latente

L_{KL} = Pérdida por divergencia

La siguiente ecuación (5) permite calcular L_{recon} , al medir la diferencia entre la imagen original (O_k) y la reconstruida (R_k), utilizando el error cuadrático medio (MSE, por sus siglas en inglés: mean squared error):

$$MSE = L_{recon} = \frac{1}{N} \sum_{k=1}^N (R_k - O_k)^2 \quad (5)$$

MSE = error cuadrático medio

R_k = Imagen reconstruida

O_k = Imagen original

Por otro lado, L_{KL} minimiza la diferencia entre la distribución latente aprendida y la distribución previa, para regularizar el espacio latente y evitar que el modelo memorice los datos de entrenamiento generando imágenes reconstruidas más precisas.

$$L_{KL} = \sum_{k=1}^N (\sigma_{rz^2} + \mu_{rz^2} - 1 - \log(\sigma_{rz^2})) \quad (6)$$

Donde:

σ_{rz^2} = desviación estándar del espacio latente

μ_{rz^2} = media del espacio latente

El parámetro β equilibra la calidad en reconstrucción de la imagen y el espacio latente. Un valor bajo β minimiza la pérdida de reconstrucción, proporcionando mayor calidad de la imagen reconstruida, pero con representaciones latentes menos estructuradas. Por el contrario, un valor alto β genera representaciones de características más consistentes e invariantes, con una menor precisión en la reconstrucción. Finalmente, el autoencoder es

utilizado para reconstruir imágenes al aprender características de las imágenes distorsionadas durante el entrenamiento. La función de pérdida, basada en el MSE, fue usada para que la red neuronal aprendiera patrones de las imágenes para su reconstrucción.

Generación del código de usuario

Generación de la matriz de características

Una vez entrenada la red neuronal siamesa, se generó el código de usuario. La construcción de este código de usuario combina una matriz de características, obtenida a partir de las salidas de las ramas de la red neuronal siamesa con la marca de agua binaria.

La matriz de características (f_v) se obtuvo mediante la siguiente fórmula (7):

$$f_v = \begin{bmatrix} V_{out1} \\ V_{out2} \end{bmatrix} \quad (7)$$

Donde:

f_v = Matriz de características

V_{out1} = Vector de salida 1 generado por la red siamesa

V_{out2} = Vector de salida 2 generado por la red siamesa

Los vectores de salida V_{out1} y V_{out2} representan las características únicas de la imagen obtenidas por las ramas de la red neuronal B_1 y B_2 y se obtienen mediante las siguientes ecuaciones (8):

$$V_{out1} = B_1(O_k) \text{ y } V_{out2} = B_2(R_k) \quad (8)$$

Donde:

V_{out1} = Vector de salida 1 generado por la red siamesa

B_1 = Rama 1 de la red neuronal

O_k = Imagen original

V_{out2} = Vector de salida 2 generado por la red siamesa

B_2 = Rama 2 de la red neuronal

R_k = Imagen reconstruida

Posteriormente, f_v se redimensionó en una matriz f_{vr} al mismo tamaño que la marca de

agua (512 x 512 píxeles), mediante escalado por vecinos más cercanos para expandir el número de elementos (9):

$$f_{vr} = \begin{bmatrix} c_{1,1} & \cdots & c_{1,n} \\ \vdots & \ddots & \vdots \\ c_{m,1} & \cdots & c_{m,n} \end{bmatrix} \quad (9)$$

Finalmente, f_{vr} se binariza (f_m) para crear el código de usuario (10):

$$f_m(i,j) = \begin{cases} 0 & \text{si } f_{vr}(i,j) \leq 0 \\ 1 & \text{si } f_{vr}(i,j) > 0 \end{cases} \quad (10)$$

La matriz de características incluye características espaciales como frecuenciales invariantes, lo que aumenta la robustez ante distorsiones o manipulaciones.

Código de usuario

El código de usuario es un elemento externo diseñado para autenticar la propiedad de la imagen de manera segura. Debe almacenarse en un dispositivo dedicado para preservar la integridad y autenticidad de la imagen original. En un dispositivo dedicado para preservar la integridad y autenticidad de la imagen original. Se obtiene utilizando la operación lógica XOR (por sus siglas en inglés: O exclusivo) y que se representa con el símbolo ($\oplus$), en la ecuación (11):

$$uc = f_m \oplus w \quad (11)$$

Donde:

uc = código de usuario uc combina lógicamente

$\oplus$ = operación lógica XOR

f_m = matriz de características binarizada extraída de la imagen distorsionada mediante la red neuronal previamente entrenada.

w = código QR

Este proceso genera una imagen binaria que parece ser ruido aleatorio. La operación XOR proporciona un elemento robusto para la autenticación de la imagen. Este método mejora la seguridad del código de usuario para evitar la autenticación no autorizada de la imagen.

Reconstrucción de imagen

Recuperación de la marca de agua para la autenticación de imágenes

El proceso de recuperación de la marca de agua permite autenticar al usuario de la imagen. Para recuperar la marca de agua, se aplicó la operación lógica XOR en la ecuación (12):

$$w' = f_m \oplus uc \quad (12)$$

Donde:

w' = marca de agua

$\oplus$ = operación lógica XOR

f_m = matriz de características binarizada extraída de la imagen distorsionada mediante la red neuronal previamente entrenada.

uc = código de usuario

La red neuronal entrenada está diseñada para generar una matriz de características similar para la misma imagen debido a la función de pérdida, incluso si ha sido distorsionada. Como resultado, la operación XOR permite recuperar una marca de agua w' que es similar a la original. La propiedad de la imagen se autenticó correctamente con la reconstrucción de la marca de agua. Una vez recuperada la marca de agua, se empleó el autoencoder entrenado para reconstruir una versión en escala de grises y a escala reducida (128×128 píxeles) de la imagen, con el propósito de verificar su autenticidad. Este proceso de reconstrucción se diseñó para restaurar imágenes que hayan sido sometidas a los tipos de distorsiones incluidas durante el entrenamiento, como traslaciones, recortes, empalmes o ataques de copia y pegado. Sin embargo, si la distorsión es severa, la imagen reconstruida puede presentar errores o, en algunos casos, resultar en una imagen completamente diferente. Además, el proceso de reconstrucción se limitó a imágenes pertenecientes al conjunto de datos utilizado durante el entrenamiento, lo que garantizó que solo reconstruyera imágenes que había aprendido a reconocer. Esta restricción disminuyó el riesgo de reconstrucciones erróneas provenientes de imágenes no relacionadas.

Evaluación de la eficacia del algoritmo

Los experimentos se realizaron utilizando las 110 imágenes sin manipular del conjunto de imágenes MICC-F220, considerando la totalidad de los resultados obtenidos, sin excluir ningún caso. Además, los experimentos se repitieron en tres ocasiones independientes, con el objetivo de validar la consistencia y reproducibilidad de los resultados. El algoritmo propuesto se evaluó mediante distintas métricas para identificar las diferencias entre las marcas de agua recuperadas y las originales: error promedio de bits (BER, por sus siglas en inglés: Bit Error Rate) recuperados, correlación cruzada normalizada (NC, por sus siglas en inglés: Normalized Cross-Correlation) y la proporción máxima de señal a ruido de la imagen (PSNR, por sus siglas en inglés: Peak Signal-to-Noise Ratio). Estas métricas ofrecen una evaluación cuantitativa de la calidad y fidelidad de la marca de agua recuperada, reflejando su precisión y resistencia ante distorsiones.

El BER mide el porcentaje de píxeles incorrectos en comparación con la marca de agua original, evaluando la precisión y los errores en el proceso de recuperación (13).

$$BER = \frac{\text{Píxeles incorrectos}}{\text{Píxeles totales}} \quad (13)$$

La NC analiza la similitud entre la marca de agua recuperada w' y la original w , mediante la correlación entre sus píxeles. Un valor alto indica gran similitud, mientras que uno bajo refleja diferencias significativas.

$$NC = \frac{\sum (w'(x, y) - \mu_{w'})(w(x, y) - \mu_w)}{\sum \sqrt{(w'(x, y) - \mu_{w'})^2 (w(x, y) - \mu_w)^2}} \quad (14)$$

Donde:

w' = marca de agua recuperada

w = marca de agua original

μ_w y $\mu_{w'}$ = los valores promedio de w y w'

El PSNR cuantifica la relación entre la imagen original en escala de grises y la imagen reconstruida, evaluando la calidad visual. Un

valor alto de indica menor distorsión y mejor calidad de recuperación (15).

$$PSNR = 10 \log_{10} \left(\frac{\max^2}{MSE} \right) \quad (15)$$

Donde:

max = valor máximo de los pixeles

MSE = Error cuadrático medio

Las pruebas experimentales incluyeron diversas distorsiones geométricas y de procesamiento de señal para validar la robustez del algoritmo.

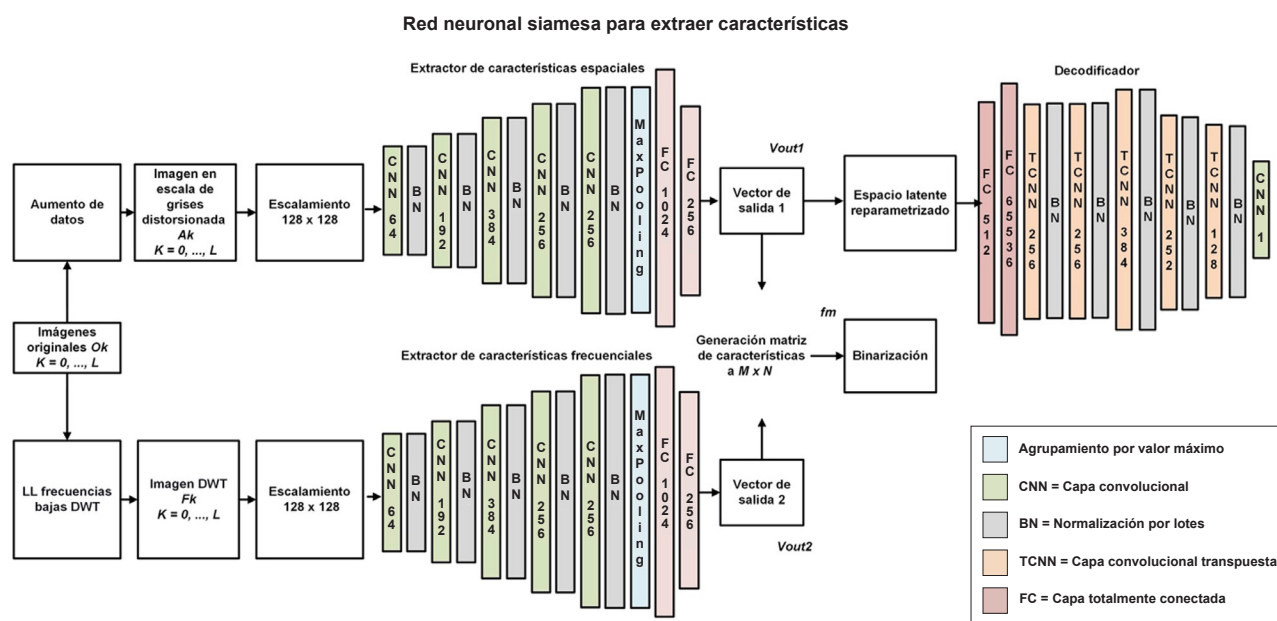
RESULTADOS Y DISCUSIÓN

Eficacia del proceso de recuperación de la marca de agua

El principal aporte del método propuesto, con una red neuronal siamesa basada en el modelo AlexNet (Figura 1), radica en su enfoque orientado a la identificación del propietario de una imagen, mediante la extracción y combinación de características robustas espaciales (ruta superior en la Figura 1) y frecuenciales (ruta inferior), obtenidas por la red neuronal. En el proceso se capturan simul-

táneamente patrones consistentes e invariantes presentes en la representación frecuencial, a través de la DWT, así como en la representación espacial en escala de grises. Mediante su posterior integración (parte central del esquema), es posible detectar alteraciones en la estructura de la imagen, lo que facilita su reconstrucción sin distorsiones significativas. Por otra parte, optimizar la eficiencia computacional permitió reducir los tiempos de entrenamiento, sin comprometer la precisión, asimismo, el desempeño del autoconder mostró ser eficiente para la reconstrucción de imágenes.

A diferencia de enfoques previos, como el de Arevalo-Ancona y col. (2024), que emplean redes siamesas centradas en características obtenidas de una red convolucional (FCN, por sus siglas en inglés: Fully Convolutional Networks) para autenticar al usuario, el método propuesto incorpora un Autoencoder Variacional (VAE, por sus siglas en inglés: Variational Autoencoders) en su arquitectura. En este diseño, la salida de la rama encargada para la extracción de características de la



■ Figura 1. Arquitectura de la red neuronal: combinación de red siamesa para la extracción de características y reconstrucción de la imagen manipulada mediante un decodificador.

Figure 1. Neural Network Architecture: Combining a Siamese Network for Characteristics Extraction and reconstruction of the manipulated image by means of a decoder.

imagen en escala de grises se utiliza como entrada del decodificador del VAE, lo que permite reconstruir la imagen original a partir de su representación latente. Esta capacidad de reconstrucción no solo mejora la verificación de la propiedad, sino que también proporciona una herramienta eficaz para la detección de manipulaciones y la restauración de imágenes alteradas.

Los parámetros de arquitectura de la red neuronal propuesta determinan su capacidad de aprendizaje y desempeño (Tabla 1). En la Figura 2a se muestra el diagrama operacional, mediante el cual, la red neuronal ex-

trae las características de la imagen original (compuestas por los patrones consistentes e invariantes identificados en las representaciones frecuenciales y espaciales de la misma) y se les integra la información del usuario para generar el código de usuario. En la Figura 2b se ilustra el proceso de recuperación de la marca de agua (cuando esta ha sido alterada de forma no autorizada), que consiste en reconocer los elementos de la imagen distorsionada e incorporarles el código de usuario original (del propietario de la imagen) para recuperar las características de la imagen y establecer en ella los datos de autenticidad original.

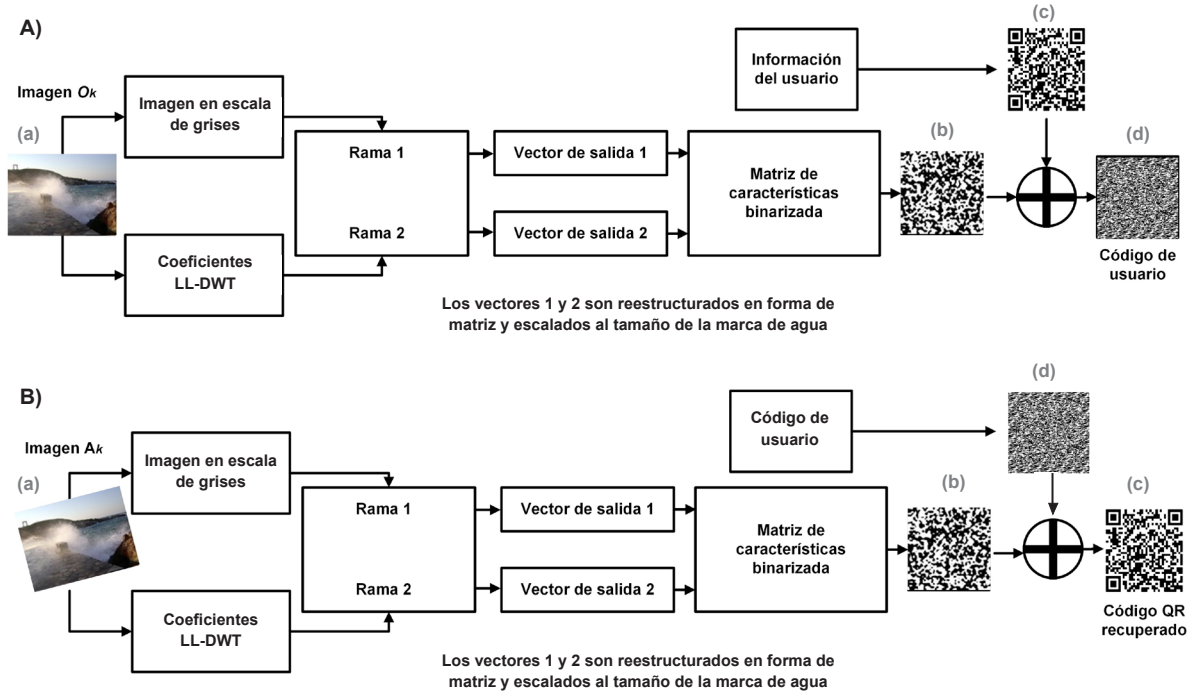
■ **Tabla 1. Configuración de las redes neuronales.**

Table 1. Neural networks configuration.

Capa	Neuronas	Tamaño del filtro	Paso	Relleno del mapa de características	Función de activación
Red neuronal siamesa para extraer características y el vector del espacio latente (entrada del decodificador)					
CNN 64	64	3 x 3	1	1	ReLu
CNN 192	192	3 x 3	2	1	ReLu
CNN 384	384	3 x 3	2	1	ReLu
CNN 256	256	3 x 3	2	1	ReLu
CNN 256	256	3 x 3	2	1	ReLu
FC 1024	1 024	---	---	---	ReLu
FC 256 (Vectores de salida)	256	---	---	---	ReLu
Decodificador del autoencoder para reconstruir imágenes					
FC 512	512	---	---	---	ReLu
FC 65536 (Redimensión [256, 16, 16])	65 536	---	---	---	ReLu
TCNN 256	256	3 x 3	1	1	ReLu
TCNN 256	256	3 x 3	2	1	ReLu
TCNN 384	384	3 x 3	2	1	ReLu
TCNN 252	252	3 x 3	2	1	ReLu
TCNN 128	128	3 x 3	2	1	ReLu
TCNN 1	1	3 x 3	2	1	Sigmoide
Hiperparámetros para la generación del código de usuario			Hiperparámetros para la reconstrucción de la imagen		
Hiperparámetros: Épocas = 20; Tamaño del lote de imágenes = 64; Tasa de aprendizaje = 0.001; Momento = 0.9; Optimizador = Adam			Hiperparámetros: Épocas = 15 000; Tamaño del lote de imágenes = 64; Tasa de aprendizaje = 0.000 05; Momento = 0.9; Optimizador = Adam		

El método propuesto mejoró la seguridad del código de usuario para evitar la autenticación no autorizada de la imagen (Figura 3). La Figura 4 presenta un ejemplo de la marca de agua recuperada de una imagen distorsionada. Las imágenes de los códigos QR contienen

información útil para evaluar visualmente la recuperación (Figura 5). La técnica obtenida en este estudio mejoró la extracción de características e identificó características robustas de la imagen, incluso en presencia de distorsiones (Chakraborty y col., 2024). Estas



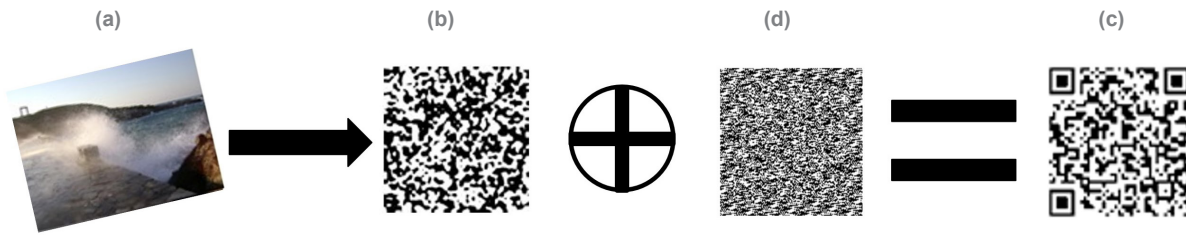
■ Figura 2. Diagrama de generación del código de usuario; A) Generación del código de usuario (resultado de un proceso de encriptación, que combina el código QR original con las características de la imagen [extraídas por la red neuronal a partir de la imagen] a través de una operación lógica XOR: Ecuación 11); B) Recuperación de la marca de agua (Código QR recuperado, obtenido a partir de una imagen [posiblemente manipulada] y sus características extraídas, con el fin de verificar su autenticidad: Ecuación 12).

Figure 2. User code generation diagram; A) Generation of the user code (result of an encryption process, which combines the original QR code with the characteristics of the image [extracted by the neural network from the image] through an XOR logical operation: Equation 11); B) Recovery of the watermark (QR code recovered, obtained from an image [possibly manipulated] and its extracted characteristics, in order to verify its authenticity: Equation 12).



■ Figura 3. Proceso de generación de código de usuario; (2A) (a) Imagen original, (b) Matriz de características, (c) Código QR, (d) Código de usuario.

Figure 3. User code generation process; (2A) (a) Original image, (b) Features matrix, (c) QR code, (d) User code.



■ Figura 4. Proceso de recuperación de la marca de agua; (2B) a) Imagen distorsionada, b) Matriz de características, d) Código de usuario, c) Código QR recuperado.

Figure 4. Watermark process recovery; (2B) (a) Distorted image, (b) Features matrix, (d) User code, (c) Retrieved QR code.

	Rotación 110° sin recorte	Rotación 20° con recorte	Traslación en y = 150 píxeles	Traslación en x = 150, y = 150 píxeles	Transformada Afin	Recorte en el centro 100 x 100	Imagen escalada 64 x 64
Imagen distorsionada							
Código QR recuperado							
	BER = 0.009 NC = 0.996	BER = 0.007 NC = 0.994	BER = 0.009 NC = 0.995	BER = 0.009 NC = 0.992	BER = 0.009 NC = 0.997	BER = 0.001 NC = 0.999	BER = 0.001 NC = 0.999
	Compresión JPEG con QF = 30	Ruido sal y pimienta $\sigma = 9\%$	Ruido gaussiano $\mu = 0, \sigma = 0.015$	Filtro de mediana núcleo = 5 x 5	Filtro promedio núcleo = 5 x 5	Corrección gamma $\gamma = 1.1$	Ecuilización de histograma
Imagen distorsionada							
Código QR recuperado							
	BER = 0 NC = 0.999	BER = 0.001 NC = 0.999	BER = 0.001 NC = 0.999	BER = 0.001 NC = 0.999	BER = 0.005 NC = 0.998	BER = 0.002 NC = 0.998	BER = 0.005 NC = 0.989
	Ruido sal y pimienta $\sigma = 15\%$ y corrección gamma $\gamma = 1.1$	JPEG Q = 50 y filtro de mediana núcleo = 5 x 5	Corrección gamma $\gamma = 0.85$ y recorte en el centro	Ruido sal y pimienta $\sigma = 5\%$ y rotación 50°	Compresión JPEG con QF = 50 y filtro de mediana núcleo = 5 x 5	Ruido sal y pimienta $\sigma = 15\%$ y corrección gamma $\gamma = 0.85$	
Imagen distorsionada							
Código QR recuperado							
	BER = 0.007 NC = 0.995	BER = 0.007 NC = 0.994	BER = 0 NC = 1	BER = 0 NC = 1	BER = 0.007 NC = 0.994	BER = 0.007 NC = 0.995	

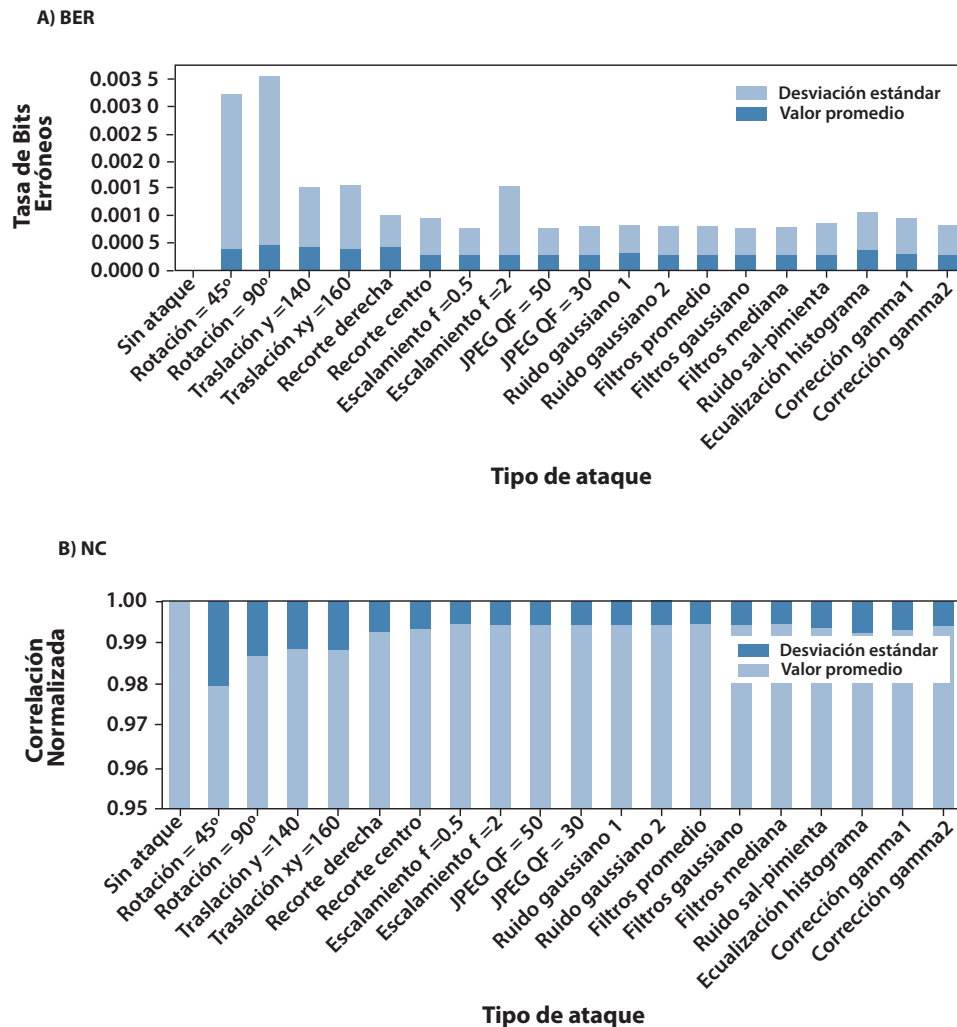
■ Figura 5. Recuperación de la marca de agua ante diferentes distorsiones. BER: Error promedio por bits; NC: Correlación cruzada normalizada.

Figure 5. Watermark recovery under different image processing distortions. BER: Bits Error Rate; NC = Normalized Cross-Correlation.

propiedades aumentaron el rendimiento y la verificación de propietario de las imágenes.

El proceso de recuperación de la marca de agua es eficiente contra la mayoría de distorsiones geométricas como rotación, traslación, transformación afín, recorte y escalado (Figura 5). Sin embargo, se obtiene un error máximo cuando una imagen específica tiene

rotación con recorte (BER = 0.007, NC = 0.994); y cuando la imagen es recortada, el valor promedio de los experimentos registró un BER de 0.002 y el NC arriba de 0.998 (Figura 6). Por otra parte, la Figura 6 muestra una desviación estándar baja, lo que indica que los datos están muy concentrados cerca de la media, es decir, que son muy similares entre sí.



■ Figura 6. Eficiencia en la recuperación de la marca de agua; A) BER (error promedio de bits), B) NC (correlación cruzada), ante distintas distorsiones (sin ataque, rotación, traslación horizontal/vertical, transformada afín, recorte central/derecho, escalamiento, JPEG QF = compresión JPEG, filtro gaussiano (núcleo 11 x 11) filtro mediana (núcleo 9 x 9), ruido sal y pimienta (densidad 0.05 %) ruido gaussiano ($\sigma = 0.09/0.2$), ecualización de histograma, corrección gamma ($y = 0.8/1.2$)).

Figure 6. Watermark Recovery Efficiency A) BER (average bit error), B) NC (cross-correlation), in the face of different distortions; (no attack, rotation, horizontal/vertical translation, affine transform, center/right cropping, scaling, JPEG QF = JPEG compression, gaussian filter (11 x 11 core) medium filter (9 x 9 core), salt and pepper noise (0.05 % density), gaussian noise ($\sigma = 0.09/0.2$), histogram equalization, gamma correction ($y = 0.8/1.2$)).

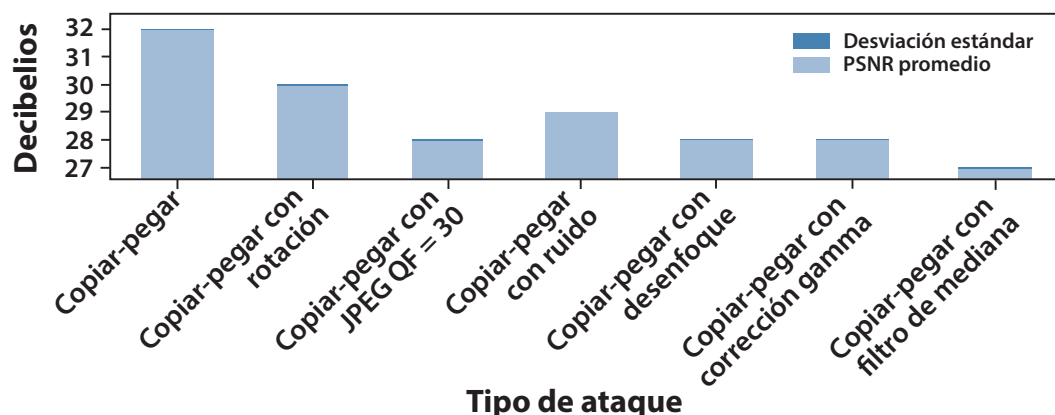
Los resultados indicaron una alta eficiencia en la recuperación de la marca de agua cuando la imagen tiene compresión JPEG (factor de calidad = 30), ruido sal y pimienta (densidad = 9 %), ruido gaussiano ($\mu = 0$, $\sigma = 0.015$) y operaciones de filtrado gaussiano y desenfoque (núcleo = 5×5). Cuando se aplica un filtro promedio, los resultados muestran un BER máximo en promedio de 0.007 y NC = 0.999.

Finalmente, también se mostró una alta robustez contra la combinación de distorsiones, logrando valores de NC superiores a 0.99 (Figura 5). La combinación de compresión JPEG y filtrado mediano o ruido sal y pimienta con corrección gamma incrementa el BER a 0.007, mientras que el NC se mantiene por encima de 0.994 (Figura 6).

El método propuesto demostró eficacia en su desempeño con técnicas existentes en términos de precisión de recuperación de la marca de agua y tiempos de procesamiento, destacando sus mejoras no solo de robustez y frente a distorsiones geométricas como rotación, traslación, transformación afín, recorte y escalado, sino también de la combinación de distintas distorsiones (Figura 7).

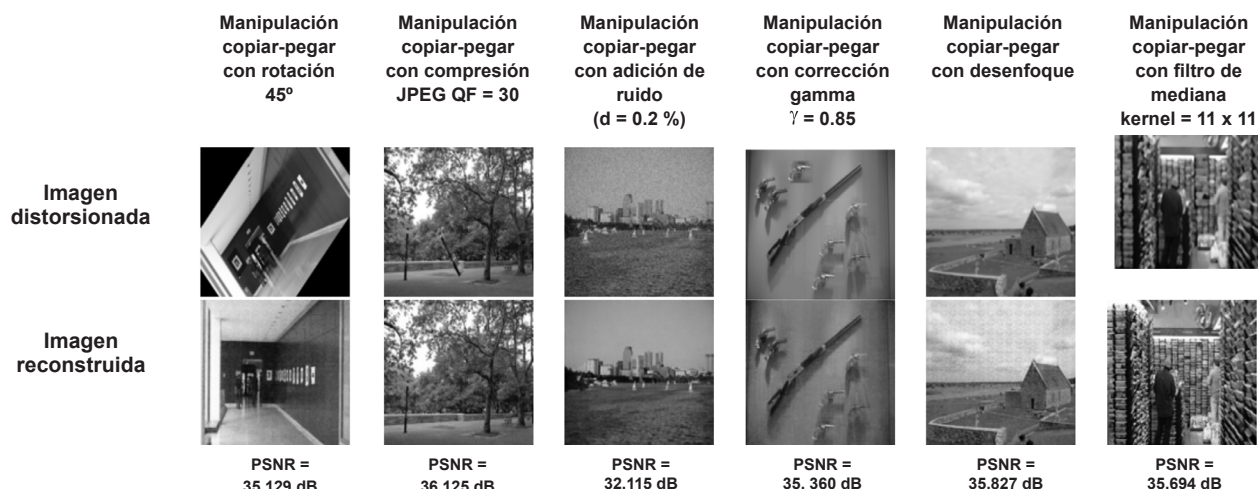
Efectividad en la reconstrucción de imágenes

La repetición experimental permitió verificar que el desempeño del método propuesto no se debe a condiciones particulares de una sola ejecución, sino que es reproducible. Aunque no se realizaron múltiples repeticiones para un análisis estadístico más amplio, estas dos ejecuciones permitieron confirmar que el comportamiento general del sistema es estable y coherente. La reconstrucción de imágenes desempeña un papel fundamental en el proceso de autenticación, ya que recupera una versión de la imagen original en escala de grises. Este procedimiento facilita la autenticación de la imagen al evidenciar posibles alteraciones o manipulaciones. La calidad en la reconstrucción de imágenes en escala de grises sometidas a diversas distorsiones y manipulaciones demostró que el método propuesto mantiene un alto PSNR, con valores promedio de los experimentos que oscilan entre 30.788 dB y 36.634 dB. La Figura 8 muestra un ejemplo particular, donde los valores oscilan entre 32.115 y 36.125. Además, el autoencoder reconstruye eficazmente imágenes alteradas por manipulaciones de tipo copia-pegar combinadas con otras distorsiones, conservando una alta fidelidad en la reconstrucción. Sin embargo, cuando las imá-



■ Figura 7. Eficiencia en la recuperación de la marca de agua evaluada mediante PSNR, que determina la relación entre la imagen original en escala de grises y la imagen reconstruida.

Figure 7. Efficiency in watermark recovery evaluated by PSNR, which determines the relationship between the original grayscale image and the reconstructed image.



■ Figura 8. Reconstrucción de la imagen en escala de grises ante diferentes distorsiones y manipulaciones.

Figure 8. Gray scale image reconstruction against distortions and manipulations.

genes son afectadas por ajustes de brillo o contraste, como la corrección gamma, la calidad de las imágenes reconstruidas disminuye, lo que evidencia una limitación del método propuesto frente a este tipo de distorsiones.

Comparación del método

El tiempo de recuperación de la marca de agua destacó la rapidez y precisión del método propuesto frente a modelos preentrenados. Tuvo una precisión mayor frente a enfoques existentes de marcado de agua libre de distorsiones, que las técnicas basadas en coeficientes de la DCT (Zhang, 2021), secuencias perceptuales (Li y col., 2023b) y modelos preentrenados como VGG o ResNet (Han y col., 2021; Darwish y col., 2023; Huang y col., 2023; Xiang y col., 2023) (Tabla 2). Asimismo, el esquema propuesto emplea una marca de agua de 512×512 píxeles, significativamente más grande que las utilizadas por métodos como Li y col. (2023b) (32×32), Darwish y col. (2023) y Hang y col. (2021; 2023) (64×64 píxeles).

Aunque modelos como VGG19, SqueezeNet y ResNet18 son eficientes en clasificación de imágenes, son menos efectivos al usarlos como extractores de características. Mientras que el método propuesto ofrece mayor precisión que otras arquitecturas neuronales preentre-

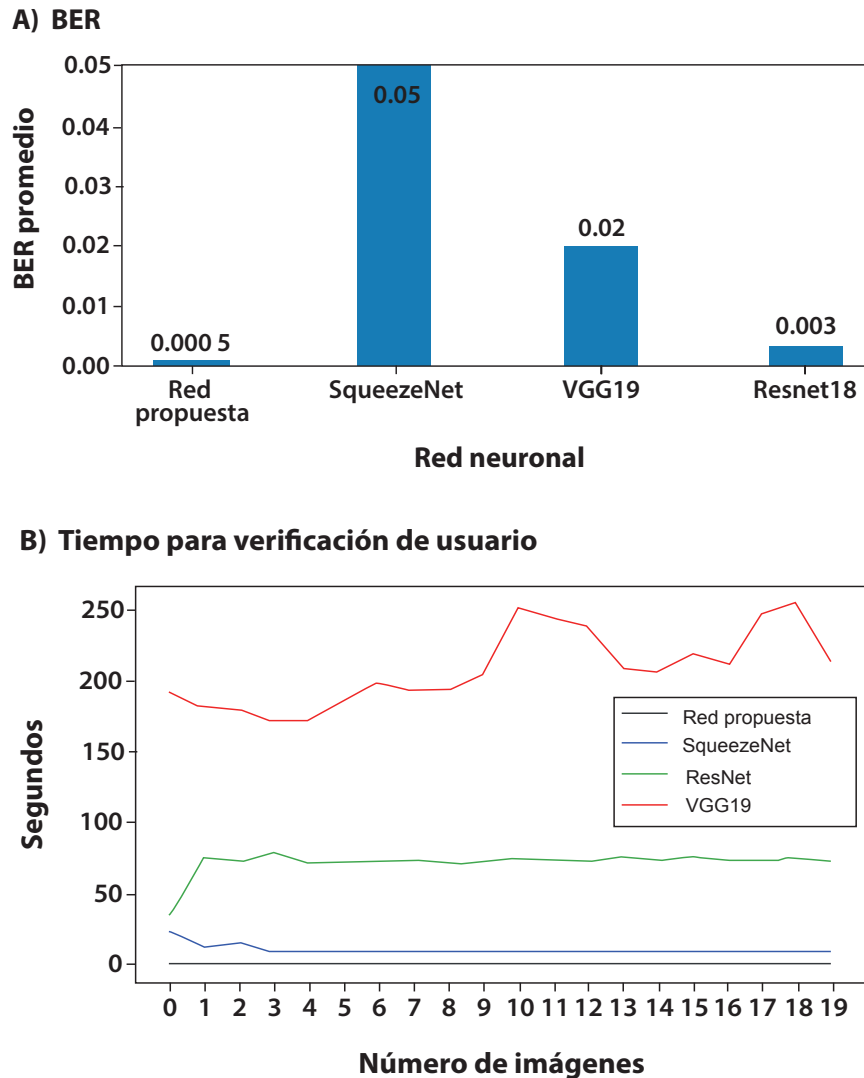
nadas, manteniendo un BER más bajo y mejor recuperación de la marca de agua (Figura 9), lo que se debe a la extracción de características mediante la red siamesa, que genera representaciones invariantes.

CONCLUSIONES

La eficacia obtenida para la autenticación de la propiedad de la imagen, por medio de la red neuronal propuesta, se debe a su diseño de doble rama. La primera rama extrae los detalles visuales en dominio espacial, y la segunda rama en el dominio DWT identifica las características generales más robustas. De esta forma, aprende al combinar la información de la imagen en su dominio espacial y de frecuencia extraídas mediante la DWT. La función de pérdida (al combinar estas características) hace que la red aprenda los detalles y las generalidades de la imagen que sean invariantes a distintos ataques. Utilizando estas características, se logró reconstruir la imagen que identifica al propietario con un valor BER menor a 0.0026 y un NC superior a 0.998. Además de la recuperación del logotipo de propiedad, se logró autenticar el contenido de la imagen original reconstruyendo su contenido con una aproximación de PSNR entre 30.788 dB y 36.634 dB. Lo anterior se obtuvo utilizando la rama espacial como codifi-

■ Tabla 2. Comparación del método propuesto.
Table 2. Proposed comparison method.

Método	Imágenes naturales				
	Método propuesto		Zhang (2021)		Darwish y col. (2023)
	Características DWT con espaciales obtenidas con red siamesa		DC coeficientes de la DCT		Características VGG16
Tamaño de la marca de agua	512 x 512		288 x 288		64 x 64
Rotación	NC = 0.999 6		NC = 0.985 5		NC = 1
Traslación	NC = 0.999 7		---		---
Recorte	NC = 0.998 9		---		---
Escalado	NC = 0.998 9		NC= 0.986 8		NC = 0.992
Filtro de mediana	NC = 0.998 3		---		---
Filtro gaussiano	NC = 0.999 3		---		NC = 0.999 2
Filtro promedio	NC = 0.999 2		---		---
Compresión JPEG	NC = 0.999 0		---		NC = 0.999 6
Adición de ruido	NC = 0.999 1		NC = 0.999 5		---
Método	Imágenes médicas				
	Li y col. (2023b)	Han y col. (2021)	Dong y col. (2023a)	Huang y col. (2023)	Xiang y col. (2023)
	Características ORB y coeficientes DCT	Características VGG19 combinado con coeficientes DFT	Características NasNet con coeficientes de la DCT	VGG sobre parametrizada	Características estadísticas de la imagen obtenidas por ResNet
Tamaño de la marca de agua	32 x 32	64 x 64	32 x 32	64 x 64	256 x 256
Rotación	NC = 0.89	NC = 0.936 2	NC = 0.95	NC = 0.982 9	NC = 0.982 9
Traslación	NC = 0.89	---	NC = 0.93	NC = 0.975 1	---
Recorte	---	NC = 0.938 6	NC = 0.86	---	NC = 0.982 9
Escalado	NC = 0.80	---	NC = 0.88	NC = 0.982 9	NC = 0.983 3
Filtro de mediana	NC = 0.80	---	NC = 0.92	NC = 0.983 3	NC = 0.983 3
Filtro gaussiano	---	---	---	NC = 0.983 3	NC = 0.983 3
Filtro promedio	---	---	---	---	---
Compresión JPEG	NC = 0.78	---	NC = 0.94	NC = 0.983 4	NC = 0.983 4
Adición de ruido	NC = 0.78	---	NC = 0.86	NC = 0.987 5	NC = 0.983 0



■ Figura 9. Comparación entre la red neuronal propuesta con redes preentrenadas; A) BER, recuperación de la marca de agua usando diferentes redes neuronales como extractores de características; B) Tiempo de procesamiento para la recuperación de la marca de agua.

Figure 9. Comparison between the proposed neural network with pretrained networks; A) BER, watermark recovery with different neural networks as features extractors; B) Watermark recovery processing time.

cador, y su correspondiente decodificador. Finalmente, para la autenticación de la propiedad, la red neuronal propuesta presentó un menor costo computacional en comparación con métodos anteriores, y una mayor variedad de ataques a los que es resistente. Como trabajo a futuro se propone ampliar el alcance del algoritmo propuesto para la autenticación de videos, evaluando su efectividad en secuencias de imágenes, adaptar el método para la detección de áreas manipuladas en imágenes y mejorar la reconstrucción en-

focándose en generar imágenes a color (RGB) y de alta definición.

AGRADECIMIENTOS

Los autores agradecen al Instituto Politécnico Nacional (IPN) y Secretaría de Ciencias, Humanidades, Tecnología e Innovación (SECIHTI) por el apoyo brindado en esta investigación.

DECLARACIÓN DE INTERESES

Los autores declararon no tener conflicto de interés alguno.

REFERENCIAS

- Amini, M., Sadreazami, H., Ahmad, M. O., & Swamy, M. N. S. (2018). A channel-dependent statistical watermark detector for color images. *IEEE Transactions on Multimedia*, 21(1), 65-73. <https://doi.org/10.1109/TMM.2018.2851447>.
- Amerini, I., Ballan, L., Caldelli, R., Del-Bimbo, A., & Serra, G. (2011). A SIFT-based forensic method for copy-move attack detection and transformation recovery. *IEEE Transactions on Information Forensics and Security*, 6(3), 1099-1110. <https://doi.org/10.1109/TIFS.2011.2129512>.
- Arevalo-Ancona, R. E., Cedillo-Hernandez, M., & Garcia-Ugalde, F. J. (2024). Robust image tampering detection and ownership authentication using zero-watermarking and Siamese neural networks. *International Journal of Advanced Computer Science and Applications*, 15(10). <https://doi.org/10.14569/IJACSA.2024.0151046>
- Aydemir, G., Paynabar, K., & Acar, B. (2022). Robust feature learning for remaining useful life estimation using Siamese neural networks. In *2022 30th European signal processing conference (EUSIPCO)* (pp. 1432-1436). IEEE. <https://doi.org/10.23919/EUSIPCO55093.2022.9909630>.
- Chakraborty, U., Thilagavathy, D., Sharma, S. K., & Singh, A. K. (2024). Hybrid deep learning with AlexNet feature extraction and U-Net classification for early detection in leaf diseases. *ICTACT Journal on Soft Computing*, 14(3), 3255-3262. <https://doi.org/10.21917/ijsc.2024.0457>
- Darwish, M. M., Farhat, A. A., & El-Gindy, T. M. (2023). Convolutional neural network and 2D logistic-adjusted-Chebyshev based zero-watermarking of color images. *Multimedia Tools and Applications*, 83(10), 29969-29995. <https://doi.org/10.1007/s11042-023-16649-3>.
- Dong, F., Li, J., Bhatti, U. A., Liu, J., Chen, Y. W., & Li, D. (2023a). Robust zero-watermarking algorithm for medical images based on improved NasNet-Mobile and DCT. *Electronics*, 12(16), 3444. <https://doi.org/10.3390/electronics12163444>.
- Dong, S., Li, J., Bhatti, U. A., Ma, J., Dong, F., & Li, Y. (2023b). Robust zero-watermarking algorithm for medical images based on GFTT-KAZE and DCT. 26th ACIS International Winter Conference on Software Engineering, Artificial Intelligence, Net-working and Parallel Distributed Computing (SNPD-Winter). <https://doi.org/10.1109/SNPD-Winter57765.2023.10223753>
- Gong, C., Liu, J., Gong, M., Li, J., Bhatti, U. A., & Ma, J. (2022). Robust medical zero-watermarking algorithm based on Residual DenseNet. *IET Biometrics*, 11(2), 135-146. <https://doi.org/10.1049/bme2.12100>.
- Han, B., Du, J., Jia, Y., & Zhu, H. (2021). Zero-watermarking algorithm for medical image based on VGG19 deep convolution neural network. *Journal of Healthcare Engineering*, (1), 5551520. <https://doi.org/10.1155/2021/5551520>.
- Huang, T., Xu, J., Tu, S., & Han, B. (2023). Robust zero-watermarking scheme based on a depth-wise overparameterized VGG network in healthcare information security. *Biomedical Signal Processing and Control*, 81, 1-10. <https://doi.org/10.1016/j.bspc.2022.104478>.
- Ikbāl, F. & Gopikakamari, R. (2022). Performance analysis of SMRT-based color image watermarking in different color. *Information Security Journal*, 31(2), 157-167. <https://doi.org/10.1080/19393555.2021.1873465>.
- Jing, W. (2020). November). Research on Digital Image Copying Watermarking Algorithm Based on Deep Learning. In *2020 International Conference on Robots & Intelligent System (ICRIS)* (pp. 104-107). IEEE. <https://doi.org/10.1109/ICRIS52159.2020.00034>.
- Khafaga, D. S., Alhammad S. M., Magdi, A., Elkomy, O., Lashin, N. A., & Hosny, K. M. (2023). Securing transmitted color images using zero-watermarking and advanced encryption standard on Raspberry Pi. *Computer Systems Science and Engineering*, 4(2), 1967-1988. <https://doi.org/10.32604/csse.2023.040345>.
- Lee, Y. H., Seo, Y. H., & Kim, D. W. (2019). Digital hologram watermarking by embedding Fresnel-diffracted watermark data. *Optical Engineering*, 58(6), 035102. <https://doi.org/10.1117/1.OE.58.3.035102>.
- Li, M., Liu, X., Wang, X., & Xiao, P. (2023a). Detecting building changes using multimodal Siamese multitask networks from very-high-resolution satellite images. *IEEE Transactions on Geoscience and Remote Sensing*, 61, 1-10. <https://doi.org/10.1109/TGRS.2023.3290817>.
- Li, Y., Li, J., Bhatti, U. A., Ma, J., Li, D., & Dong, F. (2023b). Robust zero-watermarking algorithm for medical images based on ORB and DCT. 26th ACIS

International Winter Conference on Software Engineering, Artificial Intelligence, Networking and Parallel/Distributed Computing (SNPD-Winter). <https://doi.org/10.1109/SNPD-Winter57765.2023.10223992>.

Liu, Q. L., Yang, S. Q., Liu, J., Xiong, P. C., & Zhou, M. C. (2020). A discrete wavelet transform and singular value decomposition-based digital video watermark method. *Applied Mathematical Modelling*, 85, 273-293. <https://doi.org/10.1016/j.apm.2020.04.015>.

Liu, Y. & Zhang, Z. (2021). Zero-watermarking algorithm based on DC component in DCT domain. 2021 International Conference on Electronic Information Engineering and Computer Science (EIECS). <https://doi.org/10.1109/EIECS53707.2021.9588068>.

Lin, T. Y., Maire, M., Belongie, S., Bourdev, L., Girshick, R., Hays, J., Perona, P., Ramanan, D., Zitnick, C. L., & Dollár, P. (2015). Microsoft COCO: Common objects in context. *ArXiv*. [En línea]. Disponible en: <https://doi.org/10.48550/arXiv.1405.0312>. Fecha de consulta: 15 de enero de 2025.

Shamia, D., Balasamy, K., & Suganyadevi, S. (2023). A secure framework for medical image by integrating watermarking and encryption through fuzzy-based ROI selection. *Journal of Intelligent & Fuzzy Systems*, 44, 7449-7457. <https://doi.org/10.3233/JIFS-222618>.

Shang, C., Xue, Y., Liu, W. X., & Liu, Y. (2023). Visual image digital watermarking embedding algorithm combining 3D Boolean CNN and Arnold technology. *IAENG International Journal of Computer Science*, 50(4), 1221-1231.

Solorzano, C. & Tsai, D. M. (2023). Watermark detection in CMOS image sensors using cosine-convolutional semantic networks. *IEEE Transactions on Semiconductor Manufacturing*, 36(2), 279-290. <https://doi.org/10.1109/TSM.2023.3245606>.

Vizváry, L., Sopiak, D., Oravec, M., & Bukovčíková, Z. (2019). Image quality detection using the Siamese convolutional neural network. In *2019 International Symposium ELMAR* (pp. 109-112). <https://doi.org/10.1109/ELMAR.2019.8918678>. IEEE

Wiggers, K. L., Britto, A. S., Heutte, L., Koerich, A. L., & Oliveira, L. S. (2019). Image retrieval and pattern spotting using Siamese neural network. 2019 International Joint Conference on Neural Networks (IJCNN). [En línea]. Disponible en: <https://doi.org/10.1109/IJCNN.2019.8852197>. Fecha de consulta:

10 de febrero de 2025.

Wu, J. Y., Huang, W. L., Zuo, M. J., & Gong, L. H. (2020). Optical watermark scheme based on singular value decomposition ghost imaging and particle swarm optimization algorithm. *Journal of Modern Optics*, 19(7), 1-13. <https://doi.org/10.1080/09500340.2020.1810346>.

Xiang, R., Liu, G., Li, K., Liu, J., Zhang, Z., & Dang, M. (2023). A zero-watermark scheme for medical image protection based on style feature and ResNet. *Elsevier*, 86(A), 1-11. <https://doi.org/10.1016/j.bspc.2023.105127>.

Ye, H., Huang, X., Zhu, H., & Cao, F. (2025). An enhanced network with parallel graph node diffusion and node similarity contrastive loss for hyperspectral image classification. *Digital Signal Processing*, 158, 104965. <https://doi.org/10.1016/j.dsp.2024.104965>.

Zhang, Z. (2021). Zero-watermarking algorithm based on DC component in DCT domain. 2021 International Conference on Electronic Information Engineering and Computer Science (EIECS). [En línea]. Disponible en: <https://doi.org/10.1109/EIECS53707.2021.9588068>. Fecha de consulta: 10 de enero de 2025.

Zhong, X. & Shih, F. Y. (2019). A robust image watermarking system based on deep neural networks. *ArXiv*. [En línea]. Disponible en: <https://doi.org/10.48550/arXiv.1908.11331>. Fecha de consulta: 15 de enero de 2025.