

Skema Pendanaan: MANDIRI

LAPORAN PENELITIAN



PERFORMANCE ANALYSIS OF DWT-BASED ECHO HIDING AUDIO STEGANOGRAPHY UNDER IMPERCEPTIBILITY AND ROBUSTNESS EVALUATION

TIM PENELITIAN

Ketua : Dr. Ir. Mardi Hardjianto, M.Kom. (920024)

Anggota : Dewi Kusumaningsih, M.Kom. (070024)

**FAKULTAS TEKNOLOGI INFORMASI
UNIVERSITAS BUDI LUHUR
Februari 2026**

a. Halaman Pengesahan Laporan Akhir

HALAMAN PENGESAHAN LAPORAN PENELITIAN

Judul Penelitian : PERFORMANCE ANALYSIS of DWT-BASED ECHO
HIDING AUDIO STEGANOGRAPHY UNDER
IMPERCEPTIBILITY AND ROBUSTNESS
EVALUATION

Bidang Penelitian : Cyber Security

Ketua Peneliti

- a. Nama Lengkap : Dr. Ir. Mardi Hardjianto, M.Kom.
- b. NIP/NIDN/ID-SINTA : 920024/0422036901/5977626
- c. Jabatan Fungsional : Lektor Kepala
- d. Program Studi : S2-Pasca Sarjana Magister Ilmu Komputer
- e. Nomor HP : 081586603011
- f. Alamat e-mail : mardi.hardjianto@budiluhur.ac.id

Anggota Peneliti (1)

- a. Nama Lengkap : Dewi Kusumaningsih, M.Kom.
- b. NIP/NIDN/ID-SINTA : 070024/0310128401/5992822

Mahasiswa (1)

- a. Nama Lengkap : Jeffrey Bram Pattipeilohy
- b. NIM : 2211600370

Mahasiswa (1)

- a. Nama Lengkap : Blasius Dala Nai
- b. NIM : 2311600825

Lama Penelitian : 6 bulan

Biaya Penelitian

- a. Sumber Mandiri : Rp 7.666.000,-
- b. Sumber lain (sebutkan jika ada) : Rp -

Jakarta, 27 Februari 2026

Mengetahui,
Dekan Fakultas Teknologi Informasi

Ketua Pelaksana



(Dr. Ir. Achmad Solichin, M.T.I.)
NIP 050023

(Dr. Ir. Mardi Hardjianto, M.Kom.)
NIP 920024

Menyetujui,
Direktur Riset dan Pengabdian Kepada Masyarakat

(Prof. Dr. Ir. Prudensius Maring, M.A.)

RINGKASAN

Steganografi audio merupakan teknik menyembunyikan informasi di dalam sinyal audio sambil tetap menjaga kerahasiaan pesan dan meminimalkan distorsi persepsi. Namun, upaya untuk memperoleh keseimbangan antara kualitas audio, akurasi ekstraksi, dan robustness masih menjadi tantangan, khususnya pada pendekatan yang bekerja pada domain waktu. Penelitian ini mengusulkan metode steganografi audio hibrida yang mengintegrasikan Echo Hiding dengan Discrete Wavelet Transform (DWT) dengan menyematkan informasi rahasia ke dalam koefisien detail (cD) dari sinyal audio hasil transformasi. Kebaruan penelitian ini terletak pada implementasi Echo Hiding pada koefisien detail DWT serta analisis sistematis parameter penyematan untuk mengevaluasi trade-off antara imperceptibility, akurasi ekstraksi, dan robustness pada beberapa kondisi pemrosesan sinyal yang dipilih. Eksperimen dilakukan menggunakan enam file audio berformat WAV dan dua pesan rahasia. Kinerja metode dievaluasi menggunakan Signal-to-Noise Ratio (SNR), Bit Error Rate (BER), dan Mean Opinion Score (MOS). Parameter yang dipilih, $X_0 = 60$, $X_1 = 260$, dan $\alpha = 0,7$, menghasilkan ekstraksi tanpa kesalahan dengan nilai BER = 0 dan nilai SNR berkisar dari 15,96 dB hingga 49,92 dB. Evaluasi subjektif oleh 10 responden menghasilkan nilai MOS rata-rata sebesar 4,8. Pengujian robustness menunjukkan nilai BER yang rendah pada kondisi Mild Gaussian noise, sementara kompresi MP3 pada bitrate 128 kbps menghasilkan nilai BER rata-rata sebesar 0,0561 untuk Pesan 1 dan 0,0477 untuk Pesan 2. Hasil tersebut menunjukkan adanya keseimbangan praktis antara kualitas audio dan akurasi ekstraksi pada kondisi yang dievaluasi.

PRAKATA

Puji dan syukur kami haturkan ke hadirat Tuhan Yang Maha Esa yang telah mencurahkan semua karunia sehingga diberikan keyakinan dan kemudahan dalam menyusun dan menyiapkan Laporan Penelitian ini dengan judul "PERFORMANCE ANALYSIS of DWT-BASED ECHO HIDING AUDIO STEGANOGRAPHY UNDER IMPERCEPTIBILITY AND ROBUSTNESS EVALUATION".

Pada kesempatan ini, peneliti ingin menyampaikan rasa terima kasih yang tak terhingga kepada:

1. Bapak Prof. Dr. Agus Setyo Budi, M.Sc. selaku Rektor Universitas Budi Luhur.
2. Bapak Dr. Ir. Achmad Solichin, S.Kom., M.T.I. selaku Dekan Fakultas Teknologi Informasi Universitas Budi Luhur.
3. Prof. Dr. Ir. Prudensius Maring, M.A. selaku Direktur Penelitian dan Pengabdian Kepada Masyarakat Universitas Budi Luhur.
4. Untuk semua pihak yang telah membantu penelitian ini, mohon maaf apabila ada kesalahan yang terucap dan terbersit, mohon dibukakan pintu maaf yang selapang-lapangnya.

Akhirnya dengan segala kerendahan hati, kami berharap agar penelitian ini dapat memberikan manfaat bagi peneliti, lingkungan kampus dan sekitarnya. Saran dan kritik yang membangun sangat diharapkan guna perbaikan penelitian mendatang.

Jakarta, Februari 2026
Penulis

DAFTAR ISI

RINGKASAN.....	iii
PRAKATA.....	iv
DAFTAR ISI.....	v
DAFTAR TABEL.....	vi
DAFTAR GAMBAR.....	vii
DAFTAR LAMPIRAN.....	viii
BAB I PENDAHULUAN.....	1
1. Latar Belakang.....	1
2. Pendekatan Pemecahan Masalah.....	2
3. State Of The Art dan Kebaruan.....	2
4. Peta Jalan Penelitian.....	3
BAB II METODE.....	4
1. Dataset Audio dan Pesan Rahasia.....	4
2. Algoritma Echo Hiding dan DWT.....	5
3. Proses Penyisipan Pesan (Embedding).....	5
4. Proses Ekstraksi Pesan.....	6
5. Pengujian dan Evaluasi Kinerja Sistem.....	7
BAB III HASIL PELAKSANAAN PENELITIAN.....	9
1. Data Audio.....	9
2. Proses Penyisipan Pesan.....	9
3. Proses Ekstraksi Pesan.....	11
4. Pengujian.....	11
BAB IV KESIMPULAN DAN SARAN.....	16
1. Kesimpulan.....	16
2. Saran.....	16
DAFTAR PUSTAKA.....	17
LAMPIRAN.....	19
Lampiran 1. Realisasi Penggunaan Anggaran.....	19
Lampiran 2. Catatan Harian.....	20
Lampiran 3. Artikel Ilmiah (draft/submitted/accepted/published).....	21
Lampiran 4. HKI.....	30

DAFTAR TABEL

Tabel 1 Perbandingan Metode Steganografi Audio Pada Penelitian Terdahulu.....	2
Tabel 2 Data audio (Cover Audio).....	9
Tabel 3 Pesan Untuk Penyisipan.....	12
Tabel 4 Nilai SNR, BER Stego Audio Dengan $\alpha = 0,7$	12
Tabel 5 Nilai SNR, BER Stego Audio Dengan $\alpha = 0,6$	13
Tabel 6 Perbedaan Antara Pesan Asli dan Pesan Hasil Ekstraksi Pada $\alpha = 0,6$	14
Tabel 7 Perbandingan Kinerja Metode Steganografi Audio.....	14
Tabel 8 Nilai BER Metode Yang Diusulkan Pada Serangan Derau Gaussian	14
Tabel 9 Nilai BER Metode Yang Diusulkan Pada Serangan Kompresi MP3	15

DAFTAR GAMBAR

Gambar 1 Peta Jalan Penelitian.....	3
Gambar 2. Langkah-Langkah Penelitian	4
Gambar 3 Penggalan Data Amplitudo File .WAV	9
Gambar 4 Penggalan Koefisien Detail (cD)	10
Gambar 5 Koefisien Detail (cD) yang telah disisipkan pesan rahasia.....	11
Gambar 6 Hasil Perhitungan Cepstrum	11

DAFTAR LAMPIRAN

Lampiran 1. Realisasi Penggunaan Anggaran	19
Lampiran 2. Catatan Harian.....	20
Lampiran 3. Artikel Ilmiah (draft/submitted/accepted/published)	21
Lampiran 4. HKI.....	30

BAB I

PENDAHULUAN

1. Latar Belakang

Pada era digital, pertukaran informasi sensitif dan rahasia telah menjadi bagian yang tidak terpisahkan dari komunikasi global. Transmisi data melalui internet memerlukan tingkat kerahasiaan dan keamanan yang tinggi (1,2). Upaya pengamanan juga mencakup perlindungan data dari akses tidak sah oleh pihak yang tidak bertanggung jawab (3,4). Secara umum, solusi keamanan informasi digital dapat dibedakan menjadi dua pendekatan utama, yaitu kriptografi dan penyembunyian informasi. Kedua teknologi tersebut diperlukan dalam proses penyimpanan dan transmisi data untuk menjamin keamanan informasi (5). Kriptografi merupakan komponen penting dalam berbagai sistem keamanan karena berfungsi melindungi kerahasiaan informasi sensitif (6). Proses enkripsi mengubah data menjadi bentuk yang tidak dapat dipahami oleh pihak yang tidak memiliki kunci dekripsi yang sesuai, sehingga hanya pihak yang berwenang yang dapat mengakses informasi tersebut. Perkembangan era komputasi telah menjadikan kriptografi sebagai salah satu instrumen penting dalam perlindungan data digital.

Selain kriptografi, teknik penting lainnya dalam sistem keamanan adalah penyembunyian informasi (*information hiding*). Teknik ini bekerja dengan menyisipkan informasi ke dalam suatu media. Dalam beberapa tahun terakhir, penyembunyian informasi menjadi semakin penting untuk mendukung penyimpanan dan transmisi data secara aman. Salah satu metode penyembunyian informasi digital yang paling umum digunakan adalah steganografi. Steganografi merupakan pendekatan alternatif yang bertujuan menyembunyikan keberadaan suatu pesan agar tidak menimbulkan kecurigaan dari pihak ketiga (7). Media yang dapat digunakan dalam steganografi meliputi citra, audio, dan video. Steganografi berbasis citra merupakan pendekatan yang paling banyak digunakan karena relatif mudah diimplementasikan (8). Namun, media audio menawarkan tingkat *imperceptibility* yang lebih tinggi meskipun memiliki tingkat kompleksitas yang lebih besar. Media audio digital juga memiliki redundansi yang cukup tinggi (9). Karakteristik tersebut menjadikan audio sebagai *cover object* yang potensial untuk menyisipkan data rahasia (10).

Berbagai metode telah dikembangkan dalam steganografi, antara lain *Least Significant Bit* (LSB), *Phase Coding*, *Discrete Wavelet Transform* (DWT), *Spread Spectrum*, *Discrete Cosine Transform* (DCT), serta kombinasi beberapa metode. LSB merupakan salah satu metode yang paling banyak digunakan karena menawarkan kapasitas penyisipan yang tinggi dan implementasi yang sederhana (11,12). Namun, metode LSB memiliki kelemahan karena relatif tidak tahan terhadap serangan pemrosesan sinyal, seperti kompresi, penambahan derau, dan *filtering* (13).

Salah satu teknik yang dikenal memiliki tingkat ketidakterdeteksian dan ketahanan yang baik adalah *Echo Hiding* (14). Metode ini menyisipkan data dengan menambahkan gema buatan ke dalam sinyal audio. Data dikodekan melalui variasi sejumlah parameter, seperti waktu tunda gema (*echo delay*), amplitudo, dan laju peluruhan (*decay rate*) (15). Meskipun gema dapat disamarkan di bawah ambang batas pendengaran manusia dan relatif sulit dideteksi, metode ini masih memiliki keterbatasan dalam menghadapi serangan pada domain waktu dan berbagai proses pengolahan sinyal (16).

Di sisi lain, *Discrete Wavelet Transform* (DWT) banyak digunakan dalam pemrosesan sinyal karena kemampuan merepresentasikan sinyal dalam domain waktu dan frekuensi secara simultan (17). DWT memungkinkan sinyal audio dipisahkan ke dalam beberapa subband frekuensi, sehingga proses penyisipan data dapat dilakukan pada komponen yang lebih stabil dan relatif kurang sensitif terhadap gangguan (18). Berbagai penelitian menunjukkan bahwa

metode steganografi berbasis domain transformasi, khususnya DWT, memiliki tingkat ketahanan yang lebih baik dibandingkan metode yang bekerja pada domain waktu (19–21).

Sejumlah penelitian terbaru telah mengeksplorasi berbagai teknik steganografi audio, termasuk LSB, *Echo Hiding*, dan pendekatan berbasis domain transformasi. Penelitian-penelitian tersebut pada umumnya berfokus pada peningkatan kapasitas penyisipan, *imperceptibility*, keberhasilan pemulihan pesan, atau *robustness* terhadap serangan pemrosesan sinyal. Perbandingan beberapa penelitian yang representatif disajikan pada Tabel 1.

Tabel 1 Perbandingan Metode Steganografi Audio Pada Penelitian Terdahulu

Reference	Method	Main Contribution	Remaining Issue
Imelda, Hardjianto	M-Bit LSB	Meningkatkan kapasitas penyisipan hingga empat kali lipat dengan tetap mempertahankan kualitas audio yang dapat diterima pada nilai $m \leq 4$.	Tingkat penyisipan yang lebih tinggi ($m > 4$) menimbulkan derau yang dapat terdengar, meskipun nilai SNR masih berada di atas 30 dB.
Carpentieri et al.	Echo Hiding	Mengembangkan implementasi <i>Echo Hiding</i> melalui teknik <i>frame skipping</i> , penyisipan pada kanal stereo, dan enkripsi AES dengan tetap mempertahankan gema yang tidak dapat dipersepsikan.	Distribusi gema masih memerlukan pengembangan lebih lanjut.
Rafiee & Fakhredanesh	Echo Hiding + Spread Spectrum	Meningkatkan kemampuan pemulihan pesan dan keamanan melalui modifikasi autokorelasi cepstral serta integrasi metode <i>Spread Spectrum</i> .	Kapasitas penyisipan masih dapat ditingkatkan dengan tetap mempertahankan kinerja pemulihan pesan.
Lahiri	Echo Hiding + Wavelet	Menerapkan metode <i>Echo Hiding</i> pada domain wavelet dengan menggunakan deret pseudorandom.	Analisis terhadap parameter penyisipan dan pemilihan koefisien wavelet masih terbatas.

Sebagaimana ditunjukkan pada Tabel I, penelitian-penelitian terdahulu telah membahas berbagai aspek dalam steganografi audio, meliputi kapasitas penyisipan, *imperceptibility*, pemulihan pesan, dan peningkatan keamanan. Meskipun teknik berbasis LSB memiliki kapasitas penyisipan yang tinggi dan mudah diimplementasikan, metode tersebut pada umumnya rentan terhadap berbagai proses pengolahan sinyal. *Echo Hiding* menawarkan tingkat *imperceptibility* yang lebih baik melalui penyisipan informasi dalam pola gema buatan, namun keandalan proses ekstraksinya dapat dipengaruhi oleh perubahan pada sinyal audio. Sementara itu, pendekatan berbasis DWT menyediakan kemampuan lokalisasi waktu–frekuensi dan memungkinkan proses penyisipan dilakukan pada subband frekuensi tertentu. Namun, penelitian terdahulu belum mengkaji secara memadai pengaruh parameter gema terhadap akurasi ekstraksi dan kualitas audio ketika *Echo Hiding* diterapkan secara langsung pada koefisien detail DWT.

2. Pendekatan Pemecahan Masalah

Berdasarkan permasalahan tersebut, penelitian ini mengusulkan metode steganografi audio yang menggabungkan *Echo Hiding* dan *Discrete Wavelet Transform* (DWT). Tujuan penggabungan dua metode ini untuk memanfaatkan keunggulan *Echo Hiding* dalam menjaga ketidakterdeteksian serta kemampuan DWT dalam meningkatkan ketahanan terhadap serangan pemrosesan sinyal. Proses penyisipan dilakukan dengan menerapkan *Echo Hiding* pada subband tertentu yang dihasilkan dari dekomposisi DWT, sehingga pesan tersembunyi lebih tahan terhadap noise, kompresi dan modifikasi sinyal.

3. State Of The Art dan Kebaruan

Penelitian terdahulu menunjukkan bahwa steganografi audio telah dikembangkan menggunakan berbagai pendekatan, seperti M-Bit LSB untuk meningkatkan kapasitas penyisipan, *Echo Hiding* untuk mempertahankan *imperceptibility*, kombinasi *Echo Hiding* dan *Spread Spectrum* untuk meningkatkan pemulihan pesan dan keamanan, serta penerapan *Echo Hiding* pada domain wavelet. Namun, penelitian sebelumnya masih menunjukkan keterbatasan pada aspek ketahanan, distribusi gema, kapasitas penyisipan, serta analisis parameter penyisipan dan pemilihan koefisien wavelet.

Kebaruan penelitian ini terletak pada penerapan *Echo Hiding* secara spesifik pada koefisien detail (*detail coefficients/cD*) hasil *Discrete Wavelet Transform* (DWT), disertai analisis sistematis terhadap parameter waktu tunda gema X_0 dan X_1 serta amplitudo gema (α). Penelitian ini juga mengevaluasi *trade-off* antara *imperceptibility*, akurasi ekstraksi pesan, dan *robustness* berdasarkan SNR, MOS, dan BER, termasuk pengujian terhadap *Gaussian noise* dan kompresi MP3.

4. Peta Jalan Penelitian

Peta jalan penelitian selama tiga tahun ditunjukkan pada Gambar 1 dan dibagi menjadi tiga fase yang saling berkesinambungan. Pada fase pertama, penelitian difokuskan pada pengembangan steganografi citra pada domain spasial menggunakan metode *m-bit*. Penelitian ini bertujuan menganalisis keseimbangan antara kapasitas penyisipan pesan dan kualitas citra hasil steganografi sebagai dasar pengembangan metode penyisipan data pada media digital.

Pada fase kedua, penelitian dikembangkan pada media audio digital dengan menerapkan metode *Least Significant Bit* (LSB). Fase ini diarahkan untuk menganalisis karakteristik penyisipan pesan pada sinyal audio, khususnya hubungan antara kapasitas penyisipan, kualitas audio, dan tingkat ketidak-terdeteksian pesan.

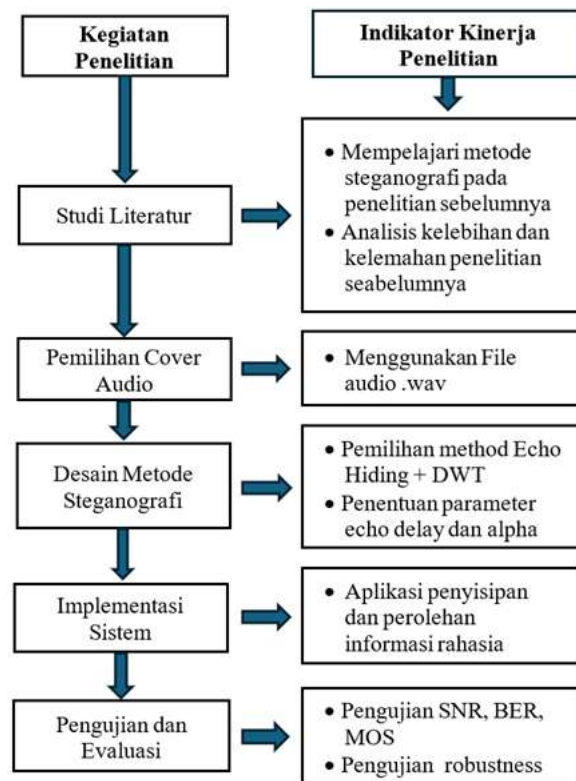
Selanjutnya, pada fase ketiga, penelitian diarahkan pada pengembangan metode steganografi yang lebih *robust* melalui pemanfaatan domain transformasi, seperti *Discrete Wavelet Transform* (DWT), serta teknik *echo hiding*. Pengembangan ini bertujuan meningkatkan ketahanan pesan tersembunyi terhadap berbagai gangguan dan manipulasi sinyal, sehingga diperoleh metode steganografi yang memiliki keseimbangan lebih baik antara kapasitas penyisipan, kualitas media, keamanan, dan *robustness*.



Gambar 1 Peta Jalan Penelitian

BAB II METODE

Penelitian ini bertujuan untuk mengevaluasi efektivitas teknik steganografi audio yang mengombinasikan *Discrete Wavelet Transform* (DWT) dan *Echo Hiding*. Kombinasi tersebut dirancang untuk menyisipkan pesan rahasia ke dalam sinyal audio digital dengan tetap mempertahankan kualitas audio secara perseptual serta mendukung proses ekstraksi pesan yang andal pada kondisi pengujian yang ditetapkan. Tahapan penelitian yang menggambarkan urutan kegiatan secara umum ditunjukkan pada Gambar 2.



Gambar 2. Langkah-Langkah Penelitian

1. Dataset Audio dan Pesan Rahasia

Dataset yang digunakan dalam penelitian ini terdiri dari audio sebagai media penampung (*cover*) dan pesan rahasia yang akan disisipkan. Audio yang digunakan berformat *.wav* dengan konfigurasi kanal mono, resolusi 16-bit serta frekuensi *sampling* 44100 Hz. Berdasarkan teorema sampling Nyquist–Shannon, suatu frekuensi dapat direproduksi secara akurat apabila frekuensi sampling yang digunakan sekurang-kurangnya dua kali lebih besar daripada frekuensi sinyal tersebut. Rentang frekuensi yang dapat didengar oleh manusia berada antara 20 Hz hingga 20 kHz, sehingga frekuensi sampling 44,1 kHz dipilih karena dinilai sesuai untuk merepresentasikan sinyal audio dalam rentang pendengaran manusia (22).

Kedalaman bit 16-bit dipilih karena format audio CD secara umum menggunakan resolusi tersebut, yang menghasilkan rentang dinamis sekitar 96 dB. Rentang dinamis sebesar 96 dB dinilai memadai untuk memperoleh kualitas perekaman audio yang baik. Format WAV dipilih karena bersifat tidak terkompresi dan memungkinkan akses langsung terhadap nilai amplitudo sinyal untuk keperluan pengolahan menggunakan metode *Discrete Wavelet Transform* (DWT) dan *Echo Hiding*. Sementara itu, pesan rahasia berupa teks (*string*) yang terlebih dahulu dikonversi ke dalam bentuk biner menggunakan standar pengodean, seperti ASCII, sebelum disisipkan ke dalam sinyal audio.

2. Algoritma Echo Hiding dan DWT

Metode penyisipan pesan yang digunakan dalam penelitian ini adalah *Echo Hiding* dan *Discrete Wavelet Transform* (DWT). Metode DWT dipilih sebagai tahap awal karena kemampuannya dalam merepresentasikan sinyal audio pada domain waktu dan frekuensi secara simultan. Pendekatan berbasis domain transformasi sering digunakan dalam proses penyisipan pesan rahasia karena dapat membantu mempertahankan tingkat *imperceptibility* serta meningkatkan ketahanan terhadap operasi pemrosesan sinyal tertentu (23).

Selain itu, penyisipan pada domain frekuensi dapat membantu mempertahankan kualitas stego-audio. Metode *Echo Hiding* dipilih sebagai teknik penyisipan karena merepresentasikan bit pesan rahasia melalui pola waktu tunda gema buatan yang tetap sulit dipersepsikan oleh sistem pendengaran manusia. Pemilihan algoritma *Echo Hiding* dan DWT dalam penelitian ini didasarkan pada kebutuhan untuk mengevaluasi *trade-off* antara kualitas audio, akurasi ekstraksi, dan *robustness* pada kondisi pengujian yang telah ditetapkan.

3. Proses Penyisipan Pesan (Embedding)

Pada tahap awal, file audio penampung (*cover audio*) dibaca dan direpresentasikan sebagai deret amplitudo dalam domain waktu. Data audio dinormalisasi ke dalam rentang -1 hingga 1. Sinyal tersebut kemudian diproses menggunakan metode Discrete Wavelet Transform (DWT) untuk memisahkan komponen frekuensi rendah dan tinggi. Komponen frekuensi rendah direpresentasikan sebagai koefisien aproksimasi (cA), sedangkan komponen frekuensi tinggi sebagai koefisien detail (cD). Perhitungan mencari nilai cA dan cD ditunjukkan pada persamaan (1).

$$\begin{aligned} cA(k) &= \sum_{n=0}^{L-1} h[n] \cdot x[2k - n] \\ cD(k) &= \sum_{n=0}^{L-1} g[n] \cdot x[2k - n] \end{aligned} \quad (1)$$

Di mana

$x[n]$ = signal audio asli (amplitude)

$h[n]$ = low-pass filter

$g[n]$ = high-pass filter

L = panjang filter (db4 = 8)

k = indeks hasil DWT

$A(k)$ = koefisien aproksimasi

$D(k)$ = koefisien detail.

Penyisipan pesan rahasia dilakukan pada koefisien detail (cD). Pemilihan cD didasari pada karakteristik sistem pendengaran manusia yang kurang sensitif terhadap perubahan kecil pada komponen frekuensi tinggi. Perubahan pada bagian ini dapat meningkatkan tingkat *interceptibility*. Pesan rahasia yang akan disisipkan terlebih dahulu diubah ke dalam bentuk biner. Setiap pesan bit disisipkan dengan menambah komponen *echo* dari dirinya sendiri dengan amplitudo kecil. Proses penyisipan ini dinyatakan melalui persamaan (2).

$$cD'(k) = cD(k) + \alpha \cdot cD(k - d) \quad (2)$$

Di mana

$cD'(k)$ = koefisien detail setelah disisipkan pada indeks k.

$cD(k)$ = koefisien detail asli pada indeks k

$cD(k - d)$ = *echo* dari koefisien detail

α = faktor penguat *echo*

k = indeks koefisien detail

d = parameter delay

Nilai amplitudo gema (α) serta parameter waktu tunda X_0 dan X_1 ditentukan melalui analisis sensitivitas awal. Sejumlah variasi nilai parameter dievaluasi untuk mengamati *trade-off* antara *imperceptibility* yang diukur menggunakan SNR dan akurasi ekstraksi yang diukur menggunakan BER. Parameter terpilih selanjutnya digunakan pada eksperimen utama untuk memperoleh keseimbangan praktis antara kualitas audio dan keandalan ekstraksi pesan.

Beberapa file *cover audio* dapat memiliki bagian awal dengan nilai koefisien detail sebesar nol. Kondisi ini dapat memengaruhi proses penyisipan karena komponen gema yang dihasilkan dari koefisien tertunda bernilai nol, juga akan bernilai nol, sehingga tidak dapat memodifikasi (cD') secara efektif. Oleh karena itu, proses penyisipan dimulai setelah ditemukan rangkaian koefisien detail tidak nol secara berurutan. Tahap ini dilakukan untuk memastikan pola gema yang disisipkan dapat terbentuk dan selanjutnya dideteksi pada proses ekstraksi.

Setelah seluruh bit pesan disisipkan ke dalam koefisien detail, sinyal audio dikembalikan ke domain waktu menggunakan *Inverse Discrete Wavelet Transform* (IDWT). Hasil proses tersebut berupa stego-audio yang secara perseptual hampir identik dengan audio asli, tetapi telah mengandung informasi tersembunyi. Pesan rahasia kemudian diperoleh kembali melalui proses ekstraksi dengan mendeteksi pola waktu tunda gema pada domain *wavelet*.

4. Proses Ekstraksi Pesan

Proses ekstraksi pesan bertujuan untuk memperoleh kembali pesan rahasia yang telah disisipkan ke dalam file audio. Ekstraksi pesan dilakukan menggunakan pendekatan *non-blind*. Parameter seperti panjang frame, nilai waktu tunda untuk bit 0 dan bit 1, serta panjang bit pesan telah diketahui sebelumnya. Oleh karena itu, proses ekstraksi memerlukan pengetahuan awal mengenai parameter-parameter tersebut, sehingga penerapan metode yang diusulkan masih terbatas pada skenario komunikasi yang tidak sepenuhnya *blind*. Pada tahap awal, file stego-audio dibaca dan dikonversi menjadi deretan nilai amplitudo dalam domain waktu. Data audio kemudian dinormalisasi ke dalam rentang -1 hingga 1 untuk menjaga kestabilan proses pengolahan sinyal. Selanjutnya, sinyal ditransformasikan menggunakan *Discrete Wavelet Transform* (DWT) sehingga menghasilkan koefisien aproksimasi (cA) dan koefisien detail (cD). Karena proses penyisipan dilakukan pada koefisien detail, proses ekstraksi juga difokuskan pada koefisien tersebut.

Koefisien detail dibagi menjadi sejumlah frame dengan panjang yang telah ditentukan. Setiap frame diasumsikan mengandung satu bit pesan rahasia. Awal frame ditentukan dengan memeriksa keberadaan rangkaian koefisien detail tidak nol secara berurutan. Untuk meningkatkan kestabilan analisis spektral dan mengurangi pengaruh diskontinuitas pada batas frame, setiap frame dikalikan dengan *Hann Window*. Tahap ini bertujuan menekan kebocoran spektral yang dapat mengganggu proses deteksi gema pada tahap berikutnya.

Selanjutnya, dilakukan perhitungan cepstrum terhadap sinyal pada setiap frame. Analisis cepstrum digunakan untuk mendeteksi keberadaan gema dengan mengidentifikasi puncak pada domain *quefreny* yang berkaitan dengan waktu tunda gema yang digunakan pada saat proses penyisipan.

Tahap berikutnya adalah mencari puncak pada rentang *quefreny* yang telah ditentukan berdasarkan nilai waktu tunda gema yang merepresentasikan bit 0 dan bit 1. Nilai tertinggi pada area cepstrum menunjukkan posisi waktu tunda gema yang paling kuat. Posisi puncak yang teridentifikasi kemudian dibandingkan dengan dua nilai waktu tunda, yaitu waktu tunda yang merepresentasikan bit 0 dan waktu tunda yang merepresentasikan bit 1.

Apabila posisi puncak lebih dekat dengan waktu tunda bit 0, frame tersebut diklasifikasikan sebagai bit 0. Sebaliknya, apabila posisi puncak lebih dekat dengan waktu tunda bit 1, frame diklasifikasikan sebagai bit 1. Hasil klasifikasi selanjutnya disimpan sebagai rangkaian bit. Rangkaian bit tersebut kemudian dikelompokkan menjadi delapan bit untuk membentuk satu byte. Setiap byte dikonversi menjadi karakter ASCII sehingga pesan teks asli dapat diperoleh kembali.

5. Pengujian dan Evaluasi Kinerja Sistem

Pengujian pada penelitian ini dilakukan untuk mengevaluasi kinerja metode steganografi yang diusulkan berdasarkan dua aspek utama, yaitu *imperceptibility*. *Imperceptibility* digunakan untuk mengukur tingkat kemiripan antara audio asli (*cover audio*) dan *audio* yang telah disisipi pesan rahasia (*stego audio*). Pengukuran dilakukan menggunakan parameter *Signal to Noise Ratio* (SNR) dengan membandingkan energi sinyal audio sebelum dan sesudah proses penyisipan. Nilai SNR dihitung menggunakan persamaan (3).

$$SNR = 10 \log_{10} \left(\frac{\sum_{i=1}^N x(i)^2}{\sum_{i=1}^N (x(i) - y(i))^2} \right) \quad (3)$$

Di mana

SNR = Rasio daya sinyal terhadap daya derau/ *noise*

$x(i)$ = sinyal audio asli

$y(i)$ = sinyal audio yang telah disisipkan

N = jumlah total sampel

Pengujian *robustness* juga dilakukan untuk mengevaluasi tingkat ketahanan pesan rahasia terhadap berbagai proses atau gangguan yang dapat terjadi pada stego-audio. Pada tahap ini, stego-audio diproses kembali untuk mengekstraksi pesan rahasia yang tersembunyi. Hasil ekstraksi kemudian dibandingkan dengan pesan asli untuk menghitung nilai *Bit Error Rate* (BER). Nilai BER dihitung menggunakan Persamaan (4).

$$BER = \frac{N_e}{N_t} \quad (4)$$

Di mana

BER = Bit Error Rate

N_e = jumlah bit yang mengalami kesalahan pada saat ekstraksi

N_t = jumlah total bit pesan yang disisipkan

Metode yang diusulkan pada dasarnya terdiri atas proses transformasi DWT/IDWT, penyisipan gema, dan ekstraksi berbasis cepstrum. Beban komputasi terutama didominasi oleh perhitungan cepstrum berbasis *Fast Fourier Transform* (FFT), sehingga kompleksitas komputasi keseluruhan diperkirakan sebesar ($O(N \log N)$). Hal ini menunjukkan bahwa metode yang diusulkan tetap layak secara komputasional untuk diterapkan pada dataset audio yang dievaluasi.

Pengujian *robustness* dilakukan dengan menerapkan beberapa serangan pemrosesan sinyal terhadap stego-audio sebelum proses ekstraksi. Dalam penelitian ini, *Additive White Gaussian Noise* (AWGN) dan kompresi MP3 pada bitrate 128 kbps digunakan sebagai skenario serangan. Pesan yang diekstraksi setelah setiap serangan kemudian dibandingkan dengan pesan asli menggunakan nilai *Bit Error Rate* (BER). Serangan berupa penyaringan (*filtering*) dan *resampling* belum disertakan dalam evaluasi penelitian ini dan dipertimbangkan sebagai bagian dari penelitian selanjutnya.

BAB III HASIL PELAKSANAAN PENELITIAN

1. Data Audio

Penelitian ini menggunakan enam file audio .wav yang digunakan sebagai *cover audio*. Jenis lagu yang digunakan vokal dan instrumental. Data audio yang digunakan ditunjukkan pada Tabel 2.

Tabel 2 Data audio (Cover Audio)

No	Nama File (.wav)	Ukuran File (Byte)	Total Sampel	Total Koefisien Detail (cD)	Daya Tampung (Karakter)
1.	In Flight	15.537.438	7768656	3884331	118
2.	Too Many Times	17.766.134	8883001	4441504	135
3.	BeautifulLiar	19.341.148	9670500	4835253	147
4.	Chinese-beat	18.081.836	9040896	4520451	137
5.	Smooth-AC	16.422.956	8211456	4105731	125
6.	Ultimatum	24.016.940	12008448	6004227	183

Jumlah total sampel diperoleh dengan membaca banyaknya sampel amplitudo pada setiap file audio. Selanjutnya, jumlah total koefisien detail (*cD*) dihitung melalui proses dekomposisi menggunakan DWT tingkat 1, yang menghasilkan jumlah koefisien detail sekitar setengah dari total sampel. Koefisien detail (*cD*) ini dimanfaatkan sebagai sarana penyisipan dalam metode *Echo Hiding*.

Kapasitas penyisipan pesan dihitung berdasarkan jumlah frame yang dapat dibentuk dari koefisien *cD*. Setiap frame terdiri atas 4.096 sampel dan digunakan untuk merepresentasikan satu bit pesan. Dengan demikian, jumlah total bit yang dapat disisipkan dihitung menggunakan $\lfloor \text{total } cD / 4096 \rfloor$. Karena satu karakter terdiri dari 8 bit, maka kapasitas dalam satuan karakter dihitung dengan $\lfloor \text{total } cD / 4096 \rfloor / 8$. Seiring dengan meningkatnya jumlah koefisien *cD*, maka kapasitas untuk menyisipkan pesan rahasia juga akan semakin besar.

2. Proses Penyisipan Pesan

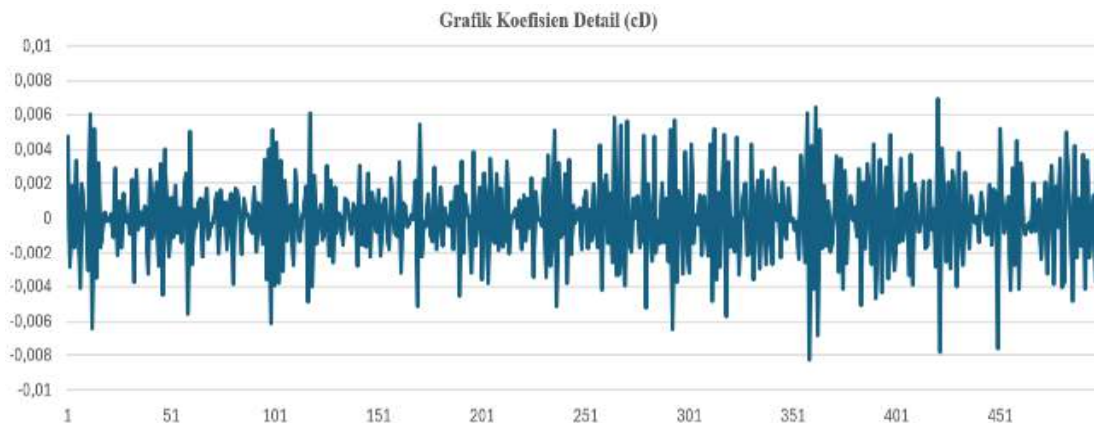
Proses penyisipan pesan dimulai dengan membaca file audio berformat .wav yang digunakan sebagai *cover audio*. Data amplitudo yang diperoleh dinormalisasi ke dalam rentang -1 hingga 1. Penggalan data amplitudo yang telah dinormalisasi ditunjukkan pada Gambar 3.



Gambar 3 Penggalan Data Amplitudo File .WAV

Sinyal audio kemudian diubah menggunakan metode *Discrete Wavelet Transform* Gelombang Diskrit (DWT) dengan wavelet Daubechies 4 (db4), sehingga membentuk komponen *cA* dan *cD*. Penyisipan pesan rahasia dilakukan pada komponen *cD* karena

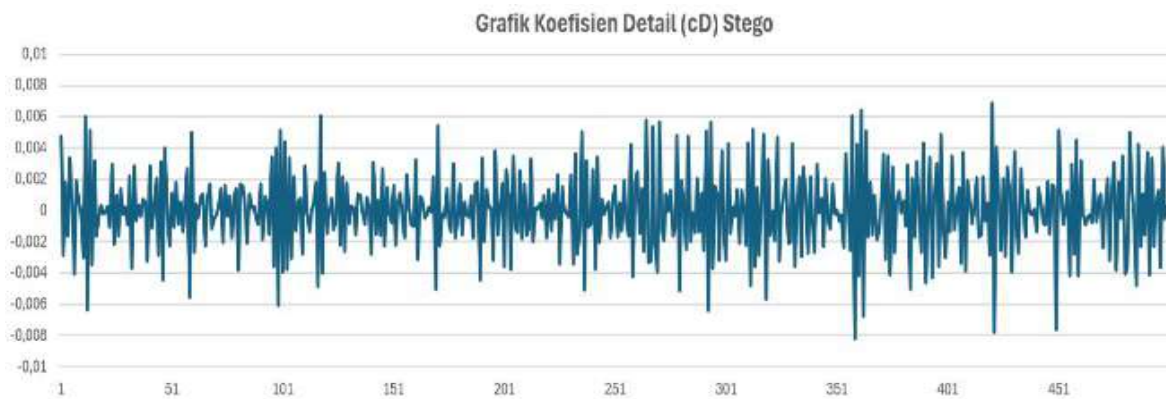
sensitivitasnya lebih rendah terhadap kemampuan pendengaran manusia. Penggalan komponen cD ditunjukkan pada Gambar 4.



Gambar 4 Penggalan Koefisien Detail (cD)

Pesan rahasia yang digunakan dalam penelitian ini berupa string. Pesan ini terlebih dahulu dikonversi ke dalam bentuk biner, kemudian disisipkan ke dalam *frame-frame* koefisien *cD*. Setiap *frame* mewakili satu bit pesan dan panjang setiap *frame* 4096 sampel. Berdasarkan hasil percobaan berulang, dipilih nilai penundaan/ delay $X_0 = 60$ untuk bit 0 dan $X_1 = 260$ untuk bit 1. Perbedaan *delay* yang cukup besar menghasilkan puncak *quefrency* yang jelas sehingga memudahkan proses ekstraksi. Nilai parameter penguatan (α) yang digunakan dalam proses penyisipan ditetapkan sebesar 0,7. Parameter ini berfungsi untuk mengatur kekuatan sinyal *echo* yang disisipkan. Pemilihan nilai ini didasarkan pada hasil percobaan agar pesan tetap dapat diekstraksi dengan baik. Salah satu pesan rahasia yang digunakan untuk melakukan penyisipan berupa teks kalimat "*The quick brown fox jumps over the lazy dog*". Setelah penyisipan pesan rahasia ke koefisien detail (cD), maka bentuk koefisien detail ditunjukkan pada Gambar 5.

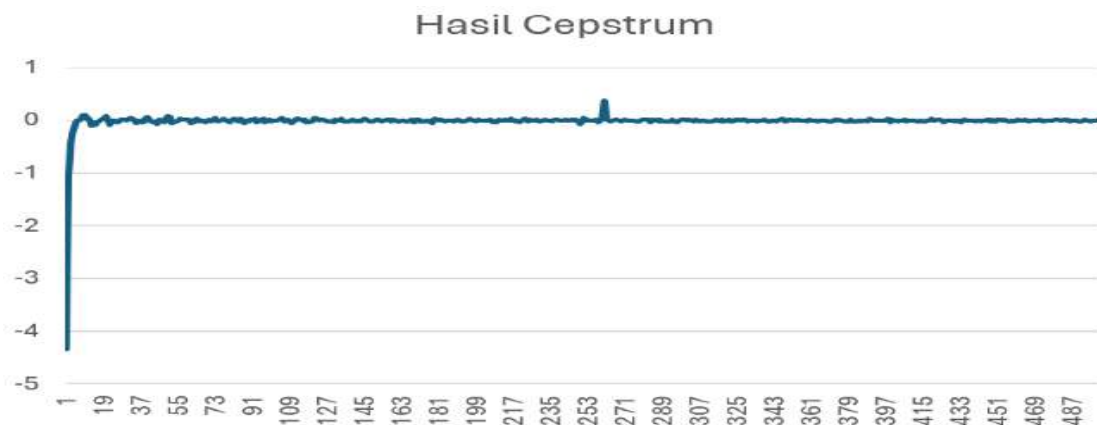
Pesan rahasia yang digunakan berupa teks (*string*). Pesan tersebut terlebih dahulu dikonversi ke dalam bentuk biner, kemudian disisipkan ke dalam *frame-frame* koefisien detail. Setiap *frame* merepresentasikan satu bit pesan dengan panjang 4.096 sampel. Berdasarkan hasil percobaan berulang, dipilih nilai waktu tunda $X_0 = 60$ untuk merepresentasikan bit 0 dan $X_1 = 260$ untuk merepresentasikan bit 1. Perbedaan nilai waktu tunda yang cukup besar menghasilkan puncak *quefrency* yang lebih jelas sehingga mempermudah proses ekstraksi. Nilai parameter penguatan gema (α) pada proses penyisipan ditetapkan sebesar 0,7. Parameter ini berfungsi untuk mengatur kekuatan sinyal gema yang disisipkan. Nilai ($\alpha = 0,7$) dipilih berdasarkan hasil eksperimen untuk menentukan ambang minimum keberhasilan ekstraksi pesan. Pengujian dengan $\alpha \leq 0,6$ tidak berhasil memperoleh kembali pesan secara sempurna, sedangkan pengujian dengan $\alpha \geq 0,7$ menunjukkan keberhasilan ekstraksi pesan. Salah satu pesan rahasia yang digunakan dalam proses penyisipan adalah kalimat "**The quick brown fox jumps over the lazy dog**", yang memuat seluruh huruf alfabet. Setelah pesan rahasia disisipkan ke dalam koefisien detail, bentuk koefisien detail hasil penyisipan ditunjukkan pada Gambar 5. Setelah seluruh bit pesan disisipkan ke dalam koefisien detail (cD), tahap selanjutnya adalah mengembalikan sinyal ke domain waktu agar dapat direpresentasikan kembali sebagai sinyal audio.



Gambar 5 Koefisien Detail (cD) yang telah disisipkan pesan rahasia

3. Proses Ekstraksi Pesan

Sinyal audio dalam format .wav dibaca dan direpresentasikan sebagai amplitudo digital. Selanjutnya dilakukan transformasi menggunakan DWT untuk memperoleh cA dan cD. Proses ekstraksi pesan dilakukan dengan menganalisis sinyal audio hasil penyisipan pada domain koefisien detail (cD). Setiap *frame* dari cD diproses menggunakan Hann Window dan dihitung cepstrum-nya untuk memperoleh representasi pada domain quefrency. Keberadaan *echo* ditunjukkan oleh puncak (peak) pada delay tertentu. Jika puncak muncul di sekitar posisi X_0 , maka merepresentasikan bit '0', sedangkan bila muncul di sekitar X_1 , maka merepresentasikan bit '1'. Dengan demikian, pesan rahasia dapat direkonstruksi kembali dari sinyal audio. Contoh hasil perhitungan cepstrum ditunjukkan pada Gambar 6. Pada hasil pengujian yang ditunjukkan pada Gambar 5, puncak utama terdeteksi berada di sekitar posisi sumbu $X \approx 261$. Posisi ini menunjukkan adanya komponen *echo* dengan delay tertentu yang telah disisipkan pada sinyal. Posisi ini lebih dekat ke X_1 dibandingkan X_0 , sehingga dapat disimpulkan bahwa *frame* tersebut merepresentasikan bit '1'.



Gambar 6 Hasil Perhitungan Cepstrum

4. Pengujian

Evaluasi kinerja dilakukan untuk menilai tingkat *imperceptibility*, akurasi ekstraksi, dan *robustness* dari metode yang diusulkan. Tingkat *imperceptibility* diukur menggunakan *Signal-to-Noise Ratio* (SNR), sedangkan akurasi ekstraksi diukur menggunakan *Bit Error Rate* (BER). Evaluasi yang dilakukan meliputi pengujian pada kondisi tanpa gangguan, analisis parameter alfa, perbandingan dengan metode pembanding, pengujian terhadap serangan *Gaussian noise*, serta pengujian terhadap kompresi MP3.

Pengujian dilakukan untuk mengevaluasi kinerja metode steganografi yang dikembangkan. Sebanyak enam file audio berformat WAV digunakan sebagai *cover audio*, sebagaimana ditunjukkan pada Tabel 2. Selain itu, dua pesan rahasia digunakan dalam proses penyisipan, sebagaimana disajikan pada Tabel 3.

Tabel 3 Pesan Untuk Penyisipan

No.	Pesan
Pesan 1	The quick brown fox jumps over the lazy dog
Pesan 2	String dengan kode ASCII mulai dari 0 hingga sebanyak daya tampung <i>cover audio</i> .

Setiap pesan disisipkan ke dalam masing-masing file *cover audio*. Setelah proses penyisipan selesai, nilai *Signal-to-Noise Ratio* (SNR) dihitung untuk mengukur tingkat *imperceptibility* stego-audio terhadap audio asli. Hasil perhitungan tersebut digunakan untuk menganalisis sejauh mana proses penyisipan pesan rahasia memengaruhi kualitas audio serta memastikan bahwa pesan dapat disisipkan tanpa menyebabkan penurunan kualitas audio yang signifikan. Hasil pengujian disajikan pada Tabel 4.

Tabel 4 Nilai SNR, BER Stego Audio Dengan $\alpha = 0,7$

Percobaan	Cover Audio	Pesan 1		Pesan 2	
		SNR (dB)	BER	SNR (dB)	BER
1.	File audio ke-1	47,11	0,00	39,69	0,00
2.	File audio ke-2	32,35	0,00	28,36	0,00
3.	File audio ke-3	24,24	0,00	19,92	0,00
4.	File audio ke-4	20,96	0,00	15,96	0,00
5.	File audio ke-5	49,92	0,00	45,45	0,00
6.	File audio ke-6	48,73	0,00	42,65	0,00

Nilai SNR yang dihasilkan dari proses penyisipan kedua pesan berkisar antara 15,96 dB hingga 49,92 dB. Meskipun nilai SNR terendah sebesar 15,96 dB, pengujian subjektif menggunakan metode *Mean Opinion Score* (MOS) yang melibatkan 10 responden menghasilkan nilai rata-rata sebesar 4,8. Hasil tersebut menunjukkan bahwa kualitas audio masih berada dalam kategori baik. Kondisi ini mengindikasikan bahwa proses penyisipan pesan tidak menimbulkan gangguan yang signifikan terhadap persepsi pendengar. Selain itu, penyisipan pada koefisien detail (*cD*) serta karakteristik gema yang terstruktur menyebabkan perubahan sinyal dipersepsikan sebagai gangguan yang tidak mengusik kualitas audio.

Pengujian dengan menyisipkan Pesan 1 ke dalam *cover audio-4* dan *cover audio-5* menghasilkan nilai SNR yang berbeda secara signifikan. Nilai SNR pada *cover audio-4* sebesar 20,96 dB, sedangkan pada *cover audio-5* sebesar 49,92 dB. Perbedaan tersebut dapat terjadi karena nilai SNR tidak hanya dipengaruhi oleh algoritma yang digunakan, tetapi juga oleh karakteristik *cover audio*. Sinyal audio dengan amplitudo yang besar cenderung menghasilkan nilai SNR yang lebih tinggi. Sebaliknya, sinyal audio dengan amplitudo rendah cenderung menghasilkan nilai SNR yang lebih rendah. Dengan demikian, nilai SNR juga dipengaruhi oleh tingkat kekuatan sinyal pada *cover audio*.

Selain karakteristik amplitudo sinyal, nilai SNR yang rendah juga dapat dipengaruhi oleh parameter α . Semakin besar nilai α , semakin kuat gema yang dihasilkan sehingga perbedaan antara sinyal asli dan sinyal stego menjadi semakin besar. Kondisi tersebut menyebabkan nilai SNR menurun. Sebaliknya, nilai α yang lebih kecil cenderung menghasilkan nilai SNR yang lebih tinggi karena perubahan pada sinyal relatif lebih kecil. Namun, peningkatan nilai α akan memperkuat pola gema yang disisipkan sehingga dapat meningkatkan tingkat keberhasilan dalam memperoleh kembali pesan rahasia.

Pada nilai $\alpha = 0,7$, seluruh pesan dapat diekstraksi dengan benar (BER = 0). Hasil ini menunjukkan bahwa kekuatan gema yang disisipkan telah memadai untuk membentuk pola waktu tunda yang jelas. Namun, ketika nilai α diturunkan menjadi 0,6, mulai terjadi kesalahan pada proses ekstraksi yang ditunjukkan oleh nilai BER > 0. Kondisi tersebut disebabkan oleh menurunnya amplitudo gema sehingga pola gema menjadi kurang dominan dibandingkan sinyal asli, terutama pada bagian sinyal yang memiliki energi rendah. Hasil ini menunjukkan bahwa penurunan nilai α dapat meningkatkan kualitas audio, tetapi pada saat yang sama menurunkan tingkat *robustness* sistem. Nilai SNR dan BER pada $\alpha = 0,6$ disajikan pada Tabel 5.

Tabel 5 Nilai SNR, BER Stego Audio Dengan $\alpha = 0,6$

Percobaan	Cover Audio	Pesan 1		Pesan 2	
		SNR (dB)	BER	SNR (dB)	BER
1.	File audio ke-1	49,34	0,00	39,7	0,00
2.	File audio ke-2	34,60	0,00	28,4	0,01
3.	File audio ke-3	26,48	0,00	19,9	0,00
4.	File audio ke-4	23,22	0,00	16,0	0,00
5.	File audio ke-5	52,03	0,02	45,5	0,01
6.	File audio ke-6	50,91	0,00	42,7	0,00

Pada *Cover Audio-5*, nilai BER untuk Pesan 1 sebesar 0,02. Nilai tersebut menunjukkan bahwa pesan yang disisipkan tidak berhasil diperoleh kembali secara sempurna. Perbedaan antara Pesan 1 dan pesan hasil ekstraksi disajikan pada Tabel 6. Kesalahan terjadi pada saat proses ekstraksi karakter spasi di antara kata “the” dan “lazy”. Karakter yang dihasilkan pada posisi tersebut berupa karakter *null*.

Tabel 6 Perbedaan Antara Pesan Asli dan Pesan Hasil Ekstraksi Pada $\alpha = 0,6$

Pesan Asli	The quick brown fox jumps over the lazy dog
Pesan setelah diekstraksi	The quick brown fox jumps over the lazy dog

Untuk memvalidasi lebih lanjut efektivitas metode yang diusulkan, dilakukan evaluasi perbandingan dengan beberapa teknik steganografi audio konvensional pada kondisi eksperimen yang sama. Metode yang dibandingkan meliputi *Least Significant Bit* (LSB), *Echo Hiding* tanpa kombinasi metode lain, steganografi LSB berbasis DWT, serta metode DWT dan *Echo Hiding* yang diusulkan. Seluruh metode diuji menggunakan file *cover audio*, pesan rahasia, dan metrik evaluasi yang sama. Hasil perbandingan berdasarkan nilai *Signal-to-Noise Ratio* (SNR) dan *Bit Error Rate* (BER) disajikan pada Tabel 7.

Tabel 7 Perbandingan Kinerja Metode Steganografi Audio

Method	Average SNR (dB)	Average BER
LSB	107.6763	0.000000
Echo Hiding	6.8433	0.000706
DWT + LSB	65.2592	0.102843
DWT + Echo Hiding	34.6117	0.000000

Sebagaimana ditunjukkan pada Tabel 7, metode LSB menghasilkan nilai rata-rata SNR tertinggi dan nilai BER sebesar nol pada kondisi ekstraksi tanpa gangguan. Namun, metode DWT + *Echo Hiding* yang diusulkan juga menghasilkan nilai BER sebesar nol dengan nilai rata-rata SNR yang berada pada tingkat sedang. Hasil ini menunjukkan bahwa metode yang diusulkan tidak melampaui kinerja LSB dari aspek *imperceptibility* pada kondisi tanpa serangan. Meskipun demikian, metode yang diusulkan menawarkan keseimbangan yang berbeda melalui kombinasi penyisipan berbasis gema dan modifikasi koefisien pada domain wavelet. Oleh karena itu, kontribusi penelitian ini lebih tepat dipahami sebagai implementasi *Echo Hiding* yang seimbang pada koefisien detail DWT, bukan sebagai metode yang secara universal lebih unggul dibandingkan seluruh metode konvensional.

Untuk mengevaluasi lebih lanjut tingkat *robustness* metode yang diusulkan, stego audio yang dihasilkan menggunakan parameter penyisipan terpilih, yaitu $X_0 = 60$, $X_1 = 260$, dan $\alpha = 0,7$, diberikan gangguan berupa *Additive White Gaussian Noise* (AWGN). Beberapa nilai standar deviasi derau (σ) digunakan untuk mensimulasikan distorsi saluran sebelum proses ekstraksi pesan dilakukan. Nilai *Bit Error Rate* (BER) yang dihasilkan pada setiap kondisi derau disajikan pada Tabel 8.

Tabel 8 Nilai BER Metode Yang Diusulkan Pada Serangan Derau Gaussian

Stego Audio	Standard Deviation (σ)				
	0.00001	0.00005	0.00010	0.00050	0.00100
Stego Audio-1	0.00000	0.01483	0.03072	0.17373	0.26801
Stego Audio-2	0.00318	0.00318	0.01271	0.03920	0.04661
Stego Audio-3	0.00000	0.00000	0.00000	0.01271	0.02966
Stego Audio-4	0.00530	0.00847	0.00845	0.01271	0.02436
Stego Audio-5	0.02966	0.12923	0.25212	0.40466	0.40466
Stego Audio-6	0.00212	0.01377	0.03708	0.24788	0.32309

Sebagaimana ditunjukkan pada Tabel 8, nilai BER secara umum meningkat seiring dengan bertambahnya standar deviasi derau (σ). Kondisi ini menunjukkan bahwa tingkat *Additive White Gaussian Noise* (AWGN) yang semakin tinggi secara bertahap merusak pola gema yang disisipkan dan menurunkan akurasi deteksi waktu tunda berbasis cepstrum. Pada tingkat derau rendah ($\sigma \leq 0,0001$), sebagian besar file *cover audio* masih menghasilkan nilai BER yang rendah. Hasil tersebut menunjukkan bahwa metode yang diusulkan memiliki toleransi terhadap distorsi saluran ringan. Namun, ketika intensitas derau meningkat menjadi ($\sigma = 0,0005$) dan ($\sigma = 0,001$), nilai BER mengalami peningkatan yang cukup nyata. Hal ini menunjukkan bahwa derau yang berlebihan mengganggu sinyal gema yang disisipkan dan menurunkan keandalan proses ekstraksi pesan.

Untuk mengevaluasi *robustness* metode yang diusulkan terhadap kompresi lossy, seluruh stego-audio yang dihasilkan dari enam file *cover audio* dan dua pesan rahasia dikompresi ke dalam format MP3 pada bitrate 128 kbps, kemudian dikonversi kembali ke format WAV sebelum proses ekstraksi dilakukan. Pengujian ini bertujuan untuk mengetahui apakah pola gema yang disisipkan pada koefisien detail DWT masih dapat dideteksi setelah melalui proses kompresi dan dekompresi. Nilai BER yang diperoleh pada pengujian kompresi MP3 disajikan pada Tabel 9.

Tabel 9 Nilai BER Metode Yang Usulkan Pada Serangan Kompresi MP3

Cover Audio	BER	
	Message 1	Message 2
Cover Audio-1	0.0731	0.0678
Cover Audio-2	0.0381	0.0350
Cover Audio-3	0.0350	0.0286
Cover Audio-4	0.0403	0.0297
Cover Audio-5	0.1335	0.1186
Cover Audio-6	0.0169	0.0064
Average	0.0561	0.0477

Sebagaimana ditunjukkan pada Tabel 9, kompresi MP3 menimbulkan kesalahan ekstraksi pada seluruh file stego-audio yang diuji. Nilai BER rata-rata yang diperoleh adalah 0,0561 untuk Pesan 1 dan 0,0477 untuk Pesan 2. Nilai BER terendah diperoleh pada *Cover Audio-6*, sedangkan *Cover Audio-5* menghasilkan nilai BER tertinggi untuk kedua pesan. Hasil tersebut menunjukkan bahwa metode yang diusulkan memiliki tingkat toleransi parsial terhadap kompresi MP3, tetapi belum dapat menjamin proses ekstraksi tanpa kesalahan setelah stego audio mengalami kompresi lossy. Variasi nilai BER yang diperoleh juga menunjukkan bahwa tingkat *robustness* terhadap kompresi dipengaruhi oleh karakteristik *cover audio* yang digunakan.

BAB IV KESIMPULAN DAN SARAN

1. Kesimpulan

Penelitian ini mengusulkan metode steganografi audio yang mengintegrasikan *Echo Hiding* ke dalam koefisien detail hasil *Discrete Wavelet Transform* (DWT). Metode yang diusulkan dievaluasi menggunakan enam file audio berformat WAV dan dua pesan rahasia. Hasil eksperimen menunjukkan bahwa parameter terpilih, yaitu $X_0 = 60$, $X_1 = 260$, dan $\alpha = 0,7$, mampu menghasilkan ekstraksi pesan tanpa kesalahan pada kondisi tanpa gangguan, dengan nilai BER = 0 untuk seluruh file audio dan pesan yang diuji. Nilai SNR yang diperoleh berkisar antara 15,96 dB hingga 49,92 dB, sedangkan evaluasi subjektif menghasilkan nilai MOS rata-rata sebesar 4,8. Hasil tersebut menunjukkan bahwa kualitas stego-audio tetap dapat diterima secara perseptual.

Hasil penelitian mengonfirmasi bahwa pemilihan parameter berpengaruh secara signifikan terhadap keseimbangan antara *imperceptibility* dan akurasi ekstraksi. Penurunan nilai α menjadi 0,6 meningkatkan nilai SNR pada beberapa kondisi, tetapi mulai menimbulkan kesalahan ekstraksi. Kondisi ini menunjukkan bahwa penyisipan gema yang lebih lemah dapat mengurangi distorsi, tetapi sekaligus menurunkan keandalan ekstraksi pesan. Evaluasi *robustness* terhadap AWGN dan kompresi MP3 menunjukkan bahwa metode yang diusulkan memiliki toleransi parsial terhadap serangan pemrosesan sinyal yang diuji, tetapi belum dapat menjamin ekstraksi tanpa kesalahan pada kondisi sinyal yang telah mengalami distorsi. Secara khusus, kompresi MP3 pada bitrate 128 kbps menghasilkan nilai BER rata-rata sebesar 0,0561 untuk Pesan 1 dan 0,0477 untuk Pesan 2. Secara keseluruhan, metode yang diusulkan memberikan keseimbangan praktis antara kualitas audio dan akurasi ekstraksi pada kondisi pengujian yang dilakukan. Namun, evaluasi *robustness* dalam penelitian ini masih terbatas pada serangan *Gaussian noise* dan kompresi MP3.

2. Saran

Penelitian selanjutnya perlu mencakup serangan berupa penyaringan dan *resampling*, penggunaan dataset audio yang lebih besar dan beragam, analisis statistik, serta perbandingan kinerja pada kondisi serangan yang identik. Pengembangan lebih lanjut juga dapat mempertimbangkan sinkronisasi adaptif, pengodean gema yang lebih kuat, dan mekanisme koreksi kesalahan untuk meningkatkan kemampuan pemulihan pesan pada kondisi distorsi sinyal yang lebih kuat.

DAFTAR PUSTAKA

1. Evsutin O, Melman A, Meshcheryakov R. Digital steganography and watermarking for digital images: A review of current research directions. *IEEE Access*. 2020;8:166589–611. doi:10.1109/ACCESS.2020.3022779
2. Yel MB, Nasution MKM. KEAMANAN INFORMASI DATA PRIBADI PADA MEDIA SOSIAL. *Jurnal Informatika Kaputama*. 2022 Jan;6(1):92–101.
3. Hazzaa F, Shabut AM, Ali NHM, Cirstea M. Security Scheme Enhancement for Voice over Wireless Networks. *Journal of Information Security and Applications*. 2021 May;58:102798. doi:10.1016/j.jisa.2021.102798
4. Radke SS, Mishra DS. Review of Image Security approaches: Concepts, Issues, Challenges and Applications. *Journal of University of Shanghai for Science and Technology*. 2021 Jun;23:1047–54.
5. Yahaya MusaM, Ajibola A. Cryptosystem for Secure Data Transmission using Advance Encryption Standard (AES) and Steganography. *International Journal of Scientific Research in Computer Science, Engineering and Information Technology*. 2019 Dec 15;317–22. doi:10.32628/cseit195659
6. Agrahari AK, Sheth M, Praveen N. Comprehensive Survey on Image Stegnography Using LSB With AES In-depth study of various requirements in the process of embedding, encrypting and extracting of text data. *International Journal of Applied Engineering Research [Internet]*. 2018. Available from: <http://www.ripublication.com>
7. Dewi RK, Munir R. PERBANDINGAN BERBAGAI METODE STEGANOGRAFI PADA CITRA DIGITAL. *Jurnal Informatika Polinema*. 2023 May 21;9(3):289–230. doi:10.33795/jip.v9i3.1336
8. Yanti F, Budayawan K. Implementasi Steganografi Menggunakan Metode Least Significant Bit (LSB) dalam Pengamanan Informasi pada Citra Digital. *Voteteknika (Vocational Teknik Elektronika dan Informatika)*. 2023;11(1):63. doi:10.24036/voteteknika.v11i1.121968
9. Dwi Hendrata A, Prihanto A. Analisis Kualitas Suara Stego Audio Penyisipan Informasi Tersembunyi dengan Metode Least Significant Bit. *Journal of Informatics and Computer Science [Internet]*. 2021;02. Available from: <https://www2.cs.uic.edu/~i101/SoundFiles/>.
10. Kaur N, Behal S. Audio Steganography Techniques-A Survey. Navneet Kaur *Int. Journal of Engineering Research and Applications* www.ijera.com [Internet]. 2014. Available from: www.ijera.com
11. Sibarani SAB, Munthe A, Purba R, Lubis AA. Pengamanan Citra Digital Menggunakan Kriptografi DnaDan Modified LSB. *Jurnal Teknologi Informasi dan Ilmu Komputer*. 2024 Dec 10;11(6):1233–42. doi:10.25126/jtiik.2024117666
12. Adhimah LF, Nurhafiyah I, Muntahar AA, Kristiaji F, Mustofa D. IMPLEMENTASI APLIKASI STEGANOGRAFI BERBASIS WEB MENGGUNAKAN ALGORITMA LSB DAN BPCS. *KOMPUTA: Jurnal Ilmiah Komputer dan Informatika*. 2023;12(2):100–8.

13. Rahman S, uddin J, Hussain H, Shah S, Salam A, Amin F, et al. A novel and efficient digital image steganography technique using least significant bit substitution. *Scientific Reports. Nature Research*; 2025 Dec. doi:10.1038/s41598-024-83147-3 PubMed PMID: 39747911.
14. Pangaribuan CE, Surbakti OK, Sihotang F, Tamba Y, Ndruru YD, Hasugian M. Analysis Of The Utilization Of Echo Hiding Algorithm In Hiding Secret Messages. *Journal Majelis Paspama*. 2024;2(2):102–7.
15. Mishra S, Yadav VK, Trivedi MC, Shrimali T. Audio steganography techniques: A survey. In: *Advances in Intelligent Systems and Computing*. Springer Verlag; 2018. p. 581–9. doi:10.1007/978-981-10-3773-3_56
16. Wen S, Tan H, Li J, Hu T. A robust audio watermarking method based on dual-encoder U-Net and short-time Fourier transform. *Cyber Security and Applications*. 2026 Jan 1;4. doi:10.1016/j.csa.2026.100129
17. Putu I, Putra RP, Ayu G, Giri VM, Raya J, Unud K, et al. Perancangan Sistem Steganografi Berbasis Transformasi Wavelet Diskrit Terintegrasi Algoritma Rijndael dan QR-Code. *JNATIA*. 2024;2(4).
18. Damanik RF, Ardhia Pramesthi R, Budiman G, Saidah S. Audio Watermarking Berbasis DWT-DCT-SVD-CPT Menggunakan QIM dan SS Audio Watermarking Based on DWT-DCT-SVD-CPT Using QIM and SS. 2019. 166–173 p.
19. Simanungkalit GA, Tarigan DS, Simangunsong DR. Discrete Wavelet Transform (DWT) Based Steganography Implementation. *Jurnal Teknik Indonesia [Internet]*. 2023. Available from: <https://jurnal.seaninstitute.or.id/index.php/jutip13>
20. Umam C, Muslih M. Enkripsi Data Teks Dengan AES dan Steganografi DWT. *InComTech : Jurnal Telekomunikasi dan Komputer*. 2023 Apr 30;13(1):28. doi:10.22441/incomtech.v13i1.15059
21. Syahdilan A, Prawira MA. IMPLEMENTASI ALGORITMA DISCRETE WAVELET TRANSFORM UNTUK MENAMBAHKAN INVISIBLE WATERMARKING PADA CITRA DIGITAL. *Jurnal Mahasiswa Teknik Informatika*. 2024.
22. Rafiee H, Fakhredanesh M. Presenting a Method for Improving Echo Hiding. *Journal of Computer and Knowledge Engineering*. 2019;2(1). doi:10.22067/cke.v2i1.74388
23. Lahiri S. Audio Steganography using Echo Hiding in Wavelet Domain with Pseudorandom Sequence. *Int J Comput Appl*. 2016 Apr 15;140(2):16–22. doi:10.5120/ijca2016909223

LAMPIRAN

Lampiran 1. Realisasi Penggunaan Anggaran

Dana Disetujui: Rp 7.666.000,

Jenis Pembelajaan	Komponen	Item	Kuantitas	Biaya Satuan	Total
Belanja Bahan	ATK	Pembelian headset/ audio monitoring	1	500.000	500.000
Belanja Bahan	ATK	Pembelian pulpen, kertas	1	100.000	100.000
Belanja Bahan	Bahan penelitian (habis pakai)	Materai	3	12.000	36.000
Pengumpulan Data	Honor pembantu peneliti	Honor Pengambilan data dan modifikasi audio	6	150.000	900.000
Pengumpulan Data	Konsumsi	Biaya konsumsi makan rapat (2 x 4 orang)	8	85.000	680.000
Analisis Data	Honor pengolah data	Honor analisis stego-audio	6	150.000	900.000
Sewa Peralatan	Peralatan penelitian	Sewa internet	6	150.000	900.000
Pelaporan penelitian	Honor administrasi peneliti	Administrasi laporan, keuangan, surat-menyerat	1	750.000	750.000
Pelaporan penelitian	Biaya Publikasi Sinta 3	Jurnal Bereputasi Nasional (terakreditasi SINTA 3)	1	2.500.000	2.500.000
Lainnya	Biaya pendaftaran HKI	HKI Aplikasi	1	400.000	400.000

Lampiran 2. Catatan Harian

No	Tanggal	Kegiatan
1.	22 Sept - 10 Okt 2025	Mencari jurnal dan mempelajari tentang steganografi berbasis audio dan teknik evaluasi
2.	13 - 17 Okt 2025	Pengumpulan data audio
3.	20 Okt – 14 Nov 2025	Merancang Model steganografi audio menggunakan metode Echo Hiding dan Discrete Wavelet Tranform (DWT)
4.	17 Nov – 19 Des 2025	Implementasi
5.	22 – 23 Des 2025	Pengujian dan evaluasi sistem
6.	5 Jan – 27 Feb 2026	Pembuatan laporan Penelitian

Performance Analysis of DWT-Based Echo Hiding Audio Steganography under Imperceptibility and Robustness Evaluation

Jeffrey Bram Pattipeilohy^[1], Mardi Hardjianto^{[2]*}, Blasius Dala Nai^[3]

Master of Computer Science ^{[1], [2], [3]}

Budi Luhur University

Jakarta, Indonesia

2211600370@student.budiluhur.ac.id^[1], mardi.hardjianto@budiluhur.ac.id^[2],

2311600825@student.budiluhur.ac.id^[3]

Abstract—Audio steganography conceals information within audio signals while preserving message confidentiality and minimizing perceptual distortion. However, balancing audio quality, extraction accuracy, and robustness remains challenging, particularly in conventional time-domain approaches. This study proposes a hybrid audio steganography method that integrates Echo Hiding with the Discrete Wavelet Transform (DWT) by embedding secret information into the detail coefficients (cD) of the transformed audio signal. The novelty of this work lies in implementing Echo Hiding in DWT detail coefficients and systematically analyzing embedding parameters to evaluate the trade-off between imperceptibility, extraction accuracy, and robustness under selected signal-processing conditions. Experiments were conducted using six WAV audio files and two secret messages. Performance was evaluated using Signal-to-Noise Ratio (SNR), Bit Error Rate (BER), and Mean Opinion Score (MOS). The selected parameters, $X_0 = 60$, $X_1 = 260$, and $\alpha = 0.7$, achieved error-free extraction under clean conditions, with BER = 0 and SNR values ranging from 15.96 dB to 49.92 dB. Subjective evaluation by 10 respondents produced an average MOS of 4.8. Robustness evaluation showed low BER under mild Gaussian noise, while MP3 compression at 128 kbps produced average BER values of 0.0561 and 0.0477 for Message 1 and Message 2, respectively. These findings indicate a practical balance between audio quality and extraction accuracy under the evaluated conditions.

Keywords: *Steganography, Echo Hiding, Discrete Wavelet Transform, Bit Error Rate, Signal-to-Noise Ratio*

I. INTRODUCTION

In the digital age, the exchange of sensitive and confidential information has become an integral part of global communication. Data transmission over the internet necessitates a high degree of confidentiality and security [1], [2]. Security efforts also include protecting data from unauthorized access by irresponsible parties [3], [4].

Cryptography and information concealing are the two types of digital information security solutions. In order to guarantee information security, these technologies are necessary for data storage and transmission. [5]. Cryptography is the key to various security systems, protecting the confidentiality of sensitive information [6]. The encryption prepare makes data garbled to anybody without the suitable decoding key, guaranteeing that as it were authorized people can get to the data. The advent of the computer age has made cryptography an important tool for data protection in digital technology.

Besides cryptography, another important technique in security systems is information hiding. It works by embedding information into a medium. In recent years, data covering up has gotten to be progressively critical for secure capacity and transmission. One of the most common methods of hiding digital information is steganography. Steganography is an alternative approach that aims to hide the existence of a message to avoid raising suspicion among third parties [7]. Media that can be utilized in steganography are pictures, sound and video. Image-based steganography is the most widely used medium because of its ease of implementation [8]. However, audio media offers a higher level of imperceptibility, despite its high complexity. Digital audio media also has considerable redundancy [9]. This makes audio media an ideal cover object for inserting confidential data [10].

There are several methods developed for steganography, such as Least Significant Bit (LSB), Phase Coding, Discrete Wavelet Transform (DWT), Spread Spectrum, Discrete Cosine Transform (DCT), and combinations of several methods. The most widely used method is LSB because it offers high capacity and a simple implementation [11], [12]. However, the LSB method is not

resistant to signal-processing attacks, such as compression, noise addition and filtering [13].

One technique known for its high undetectability and resilience is Echo Hiding [14]. This method inserts data by adding artificial echoes into the audio signal. The data is encoded through variations in parameters such as echo delay, amplitude, and decay rate [15]. Although the echo is masked below the threshold of human hearing and is relatively difficult to detect, it still has limitations in terms of resilience to time-domain attacks and signal processing [16].

On the other hand, the Discrete Wavelet Transform (DWT) is frequently utilized in signal processing due to its ability to represent signals in both the time and frequency

domains simultaneously [17]. DWT allows the separation of audio signals into several frequency subbands, enabling the data embedding process to be carried out on components that are more stable and less sensitive to interference [18]. Various studies have shown that transform-based steganography, especially DWT, has better resilience than time-domain methods [19], [20], [21].

Recent studies have explored various audio steganography techniques, including LSB, Echo Hiding, and transform-domain approaches. These studies have mainly focused on improving embedding capacity, imperceptibility, message recovery, or robustness against signal-processing attacks. A comparison of representative studies is presented in TABLE I.

TABLE I. COMPARISON OF AUDIO STEGANOGRAPHY METHODS IN PREVIOUS STUDIES

Reference	Method	Main Contribution	Remaining Issue
Imelda, Hardjianto	M-Bit LSB	Increased embedding capacity up to four times while maintaining acceptable audio quality for $m \leq 4$	Higher embedding levels ($m > 4$) introduce audible noise despite SNR remaining above 30 dB
Carpentieri et al.	Echo Hiding	Improved Echo Hiding implementation using frame skipping, stereo embedding, and AES encryption while maintaining imperceptible echoes	Further improvement in echo distribution is still required
Rafiee & Fakhredanesh	Echo Hiding + Spread Spectrum	Improved message recovery and security through modified cepstral autocorrelation and Spread Spectrum integration	Embedding capacity can still be further improved while maintaining recovery performance
Lahiri	Echo Hiding + Wavelet	Applied Echo Hiding in the wavelet domain using a pseudorandom sequence	Limited analysis of embedding parameters and wavelet coefficient selection

As shown in TABLE I, previous studies have addressed different aspects of audio steganography, including embedding capacity, imperceptibility, message recovery, and security improvement. Although LSB-based techniques have a high embedding capacity and are easy to deploy, they are typically susceptible to signal-processing processes. Echo Hiding provides better imperceptibility by embedding information in artificial echo patterns, but the reliability of its extraction may be affected by changes in the audio signal. Meanwhile, DWT-based approaches provide time–frequency localization and enable embedding in selected frequency subbands, but previous studies have not sufficiently examined how echo parameters affect extraction accuracy and audio quality when Echo Hiding is applied directly to DWT detail coefficients.

Based on these limitations, the research gap addressed in this study is the lack of systematic analysis of Echo Hiding parameters in the DWT detail-coefficient domain, particularly regarding the trade-off among imperceptibility, extraction accuracy, and robustness under selected signal-processing conditions. Therefore, this study proposes a DWT-based Echo Hiding audio steganography method in which secret information is embedded into the detail coefficients (cD) of the transformed audio signal. The proposed approach is designed to evaluate the effects of delay parameters and echo amplitude on Bit Error Rate (BER), Signal-to-Noise Ratio (SNR), and message recovery performance.

The main contribution of this study is the implementation and performance evaluation of Echo Hiding in DWT detail coefficients, supported by systematic parameter analysis using X_0 , X_1 , and α . This study also evaluates the proposed method under clean conditions, a baseline method comparison, a Gaussian noise attack, and an MP3 compression attack. The findings are expected to provide insight into the practical balance between imperceptibility, extraction accuracy, and robustness in DWT-based audio steganography systems.

II. METHODOLOGY

This research aims to assess the effectiveness of an audio steganography technique that combines the Discrete Wavelet Transform (DWT) and Echo Hiding. This combination is designed to embed secret messages into digital audio signals while maintaining perceptual audio quality and supporting reliable message extraction under the evaluated conditions. The research steps used to illustrate the general sequence of activities are shown in Fig. 1.

A. Audio Dataset and Secret Messages

The dataset used in this research consists of audio as a cover medium and a secret message to be inserted. The audio used is in .wav format with a mono channel configuration, 16-bit resolution, and a frame rate of 44100 Hz. The Nyquist-Shannon sampling theorem states that to reproduce a frequency accurately, the sampling rate must be at least twice

that frequency. The human ear's ability to capture sound frequencies is from 20 Hz to 20 kHz [22], so the most suitable sample rate is 44.1 kHz. The 16-bit bit depth was chosen because audio CDs originally employed 16 bits, which translates to a dynamic range (dB) of 96 dB. A dynamic range of 96 dB is more than sufficient to achieve good recording quality. The .wav file is uncompressed and provides direct access to amplitude values for signal processing using the Discrete Wavelet Transform (DWT) and Echo Hiding methods. Meanwhile, the secret message is a text (string) converted to binary using a coding standard such as ASCII, then inserted into the audio signal.

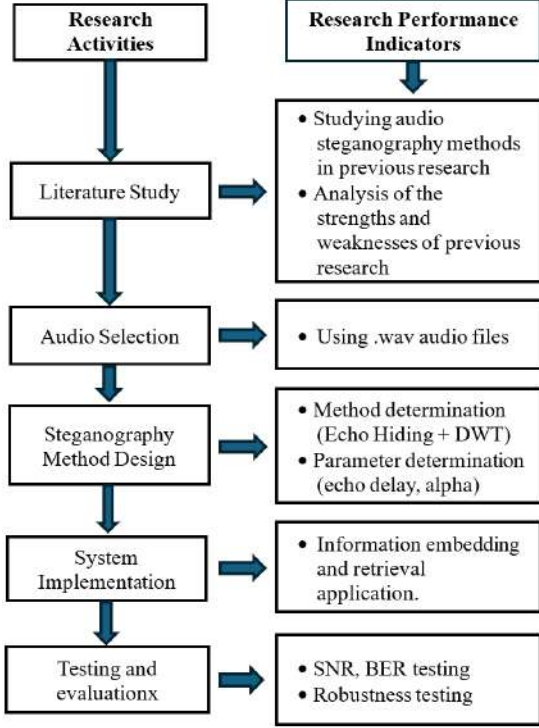


Fig. 1. Research Steps

B. Echo Hiding and DWT Algorithm

The message embedding methods chosen in this research are Echo Hiding and DWT. The DWT method was chosen as the initial stage because of its ability to represent audio signals in both the time and frequency domains. The transform-domain approach is often used for embedding secret messages because it can help maintain imperceptibility and improve resistance to certain signal-processing operations [23].

Furthermore, transform-domain embedding can help maintain the quality of stego audio. Echo Hiding was chosen as the embedding method because it represents secret bits as artificial echo-delay patterns that remain perceptually unobtrusive. The selection of the Echo Hiding and DWT algorithms in this research is based on the need to evaluate the trade-off among audio quality, extraction accuracy, and robustness under the selected testing conditions.

The Daubechies wavelet (db4) was selected for its good balance between time and frequency localization, making it suitable for non-stationary audio signals. The decomposition level was chosen to ensure sufficient separation between frequency subbands while preserving the quality of signal reconstruction. Higher decomposition levels were avoided due to potential loss of temporal resolution and increased reconstruction distortion.

C. Message Embedding Process

In the initial stage, the cover audio file is read and represented as an amplitude series in the time domain. The audio data is normalized into the range -1 to 1. The low-frequency and high-frequency components of the signal are then separated using the Discrete Wavelet Transform (DWT) technique. Approximation coefficients (cA) are used to represent low-frequency components, and detail coefficients (cD) are used to represent high-frequency components. The low-frequency components are represented as approximation coefficients (cA), while the high-frequency components are represented as detail coefficients (cD). The calculation for finding the cA and cD values is shown in (1).

$$\begin{aligned} cA(k) &= \sum_{n=0}^{L-1} h[n] \cdot x[2k - n] \\ cD(k) &= \sum_{n=0}^{L-1} g[n] \cdot x[2k - n] \end{aligned} \quad (1)$$

Where

$x[n]$ = original audio signal (amplitude)

$h[n]$ = low-pass filter

$g[n]$ = high-pass filter

L = filter length (db4 = 8)

k = DWT result index

$cA(k)$ = approximation coefficient

$cD(k)$ = detail coefficient

Secret message embedding is performed on the detail coefficient (cD). The human auditory system, which is less sensitive to little variations in high-frequency components, is the basis for the choice of cD. Small modifications to this component can improve imperceptibility, as high-frequency changes are generally less noticeable to human hearing. The secret message is transformed into binary before being embedded. Each message bit is embedded by adding an echo component of its own with a small amplitude. This embedding process is shown in (2).

$$cD'(k) = cD(k) + \alpha \cdot cD(k - d) \quad (2)$$

Where

$cD'(k)$ = detail coefficient after being embedded at index k.

$cD(k)$ = original detail coefficient at index k

$cD(k - d)$ = echo of detail coefficient

α = echo amplitude factor

k = detail coefficient index

d = delay parameter (d = X₀ for bit 0, d = X₁ for bit 1)

The values of the echo amplitude (α) and delay parameters (X₀ and X₁) were determined through a preliminary sensitivity analysis. A range of parameter values was evaluated to observe the trade-off between imperceptibility, measured by SNR, and extraction accuracy, measured by BER. The selected parameters were then used in the main experiment to achieve a practical balance between audio quality and message extraction reliability. Some cover audio files may contain initial regions with zero-valued detail coefficients. This condition can affect the embedding process because an echo component generated from a zero-valued delayed coefficient will also have a value of zero and will not effectively modify cD'. Therefore, the embedding process is initiated after a sequence of consecutive non-zero detail coefficients is found, ensuring that the embedded echo pattern can be formed and later detected during extraction.

The detail coefficients, once embedded with the message bits, then return the audio signal to the time domain using the Inverse Discrete Wavelet Transform (IDWT). The result of this process is stego audio that is perceptually almost identical to the original audio, but contains hidden information. This stego audio must be extracted to obtain the message by detecting the delay echo pattern in the wavelet domain.

D. Message Extraction Process

The message extraction process aims to recover the secret message embedded in the audio file. Message extraction is carried out using a non-blind approach. Parameters such as the frame length, the delay values for bits 0 and 1, and the message bit length are known. Consequently, the extraction process requires prior knowledge of these parameters, which limits the applicability of the proposed method in fully blind communication scenarios. In the initial stage, the stego audio file is read and converted into a sequence of time-domain amplitude values. The audio data is normalized to the range -1 to 1 to maintain stability in the signal processing. Next, the signal is transformed into the frequency domain using a DWT, producing the approximation coefficient (cA) and the detail coefficients (cD). The embedding process is performed on the detail coefficients, so the extraction process also focuses on the detail coefficients.

The detail coefficients are divided into frames of predetermined length. One bit of the secret message is assumed to be contained in each frame. The beginning of a frame is determined by checking for a sequence of consecutive nonzero detail coefficients. To improve the stability of the spectral analysis and reduce the impact of discontinuities at frame boundaries, each frame is multiplied by a Hann Window. The purpose of this step is to suppress spectral leakage that can interfere with echo detection in the

next stage. Next, the cepstrum is calculated for the signal of that frame. Cepstrum analysis is useful for detecting echoes in the signal by identifying peaks in the frequency domain associated with the echo delay used during embedding.

The next step is to find the peak within the specified quefrency range based on the echo delay values representing bit 0 and bit 1. The highest value in the cepstrum area indicates the location of the strongest echo delay. The identified peak location is then compared with two delay values, namely the delay indicating bit 0 and the delay indicating bit 1. If the peak position is closer to the delay bit 0, the frame is classified as bit 0; if it is closer to the delay bit 1, it is classified as bit 1. The classification result is then stored as a sequence of bits. This series of bits is collected and divided into eight-bit groups to form one byte. Each byte is converted into an ASCII character so that the original text message can be obtained.

E. Performance Testing and Evaluation

The research results were evaluated to assess the performance of the proposed steganography method with respect to two main aspects: imperceptibility and robustness. The degree of similarity between the original audio (cover audio) and the audio embedded with a secret message (stego audio) is measured by imperceptibility. Measurements were made using the Signal-to-Noise Ratio (SNR) by comparing the audio signal energy before and after the embedding process. The SNR value was calculated using (3).

$$SNR = 10 \log_{10} \left(\frac{\sum_{i=1}^N x(i)^2}{\sum_{i=1}^N (x(i) - y(i))^2} \right) \quad (3)$$

Where

SNR = Signal to Noise Ratio

x(i) = original audio signal

y(i) = embedded audio signal

N = total number of samples

Robustness testing is also performed to evaluate how resilient the secret message is to processes or interference that may occur in the stego audio. At this stage, the stego audio is reprocessed to extract the hidden secret message. The extraction results are then compared with the original message to calculate the Bit Error Rate (BER). The BER value is calculated using (4).

$$BER = \frac{N_e}{N_t} \quad (4)$$

Where

BER = Bit Error Rate

N_e = number of bits that are in error during extraction

N_t = total number of embedded message bits

The proposed method mainly consists of DWT/IDWT transformations, echo embedding, and cepstrum-based

extraction. The computational cost is dominated by the FFT-based cepstrum computation, resulting in an overall complexity of approximately $O(N \log N)$. This indicates that the proposed method remains computationally feasible for the evaluated audio dataset.

Robustness was evaluated by applying selected signal-processing attacks to the stego audio before extraction. In this study, additive white Gaussian noise (AWGN) and MP3 compression at 128 kbps were used as attack scenarios. The extracted message after each attack was compared with the original message using BER. Filtering and resampling attacks were not included in the current evaluation and are considered as future work.

III. RESULTS AND DISCUSSION

A. Cover Audio

The cover audio files used in this research consist of six .wav song files. The songs in the cover audio are both vocal and instrumental. The audio data used is shown in Table II.

TABLE II. COVER AUDIO

#	File Name (.wav)	File Size (Byte)	Total Sample	Total Detail Coefficient (cD)	Capacity (Character)
1.	In_Flight	15.537.438	7768656	3884331	118
2.	Too_Many_Times	17.766.134	8883001	4441504	135
3.	BeautifulLiar	19.341.148	9670500	4835253	147
4.	Chinese-beat	18.081.836	9040896	4520451	137
5.	Smooth-AC	16.422.956	8211456	4105731	125
6.	Ultimatum	24.016.940	12008448	6004227	183

The total sample is obtained by reading the number of amplitude samples in each audio file. Next, the total detail coefficients (cD) are calculated through decomposition using the DWT at level 1 and are approximately half of the total number of samples. The message capacity is calculated based on the total frames that can be formed from the cD coefficients. Each frame consists of 4096 samples, each representing one bit of the message, so the total number of bits that can be embedded is $\lfloor \text{total cD} / 4096 \rfloor$. One character consists of 8 bits, so the capacity in character units is $\lfloor \text{total cD} / 4096 \rfloor / 8$. As the number of cD coefficients increases, the capacity to embed secret messages increases.

B. Message Embedding Process

The message embedding process begins by reading the .wav audio file used as the cover audio. The resulting amplitude data is normalized to the range -1 to 1. A fragment of the normalized amplitude data is shown in Fig. 2.

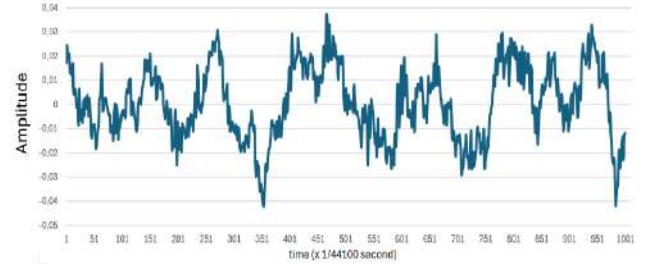


Fig. 2. Normalized .WAV File Amplitude Data Snippet

The cA and cD components are then obtained by applying the Discrete Wavelet Transform (DWT) with the Daubechies 4 (db4) wavelet to the audio signal. The secret message is inserted into the cD component because its sensitivity is lower than that of human hearing. A fragment of the cD component is shown in Fig. 3.

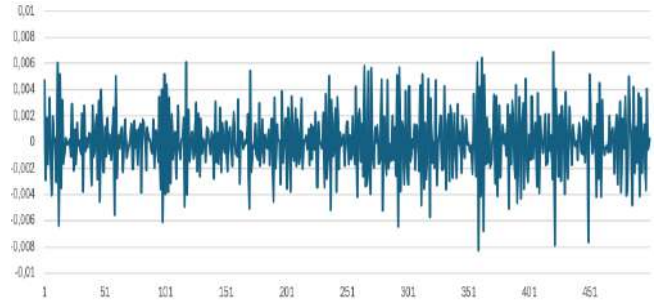


Fig. 3. Detail Coefficient (cD) Fragment

The secret message used is a string. This message is first converted into binary, then embedded into frames of detail coefficients. Each frame represents one bit of the message. The length of each frame is 4096 samples. Based on the results of repeated experiments, the delay values $X0 = 60$ for bit 0 and $X1 = 260$ for bit 1 were chosen. The large delay difference produces a clear *quefreny* peak, simplifying the extraction process. The gain parameter value (alpha) used in the embedding process is set at 0.7. This parameter controls the strength of the embedded echo signal. An alpha value of 0.7 was chosen based on experimental results to determine the minimum threshold for successful message extraction. Experiments with $\alpha \leq 0.6$ failed to retrieve messages, while $\alpha \geq 0.7$ succeeded. One of the secret messages used to embed text is the sentence "The quick brown fox jumps over the lazy dog," in which all the letters are used. After embedding the secret message into the detail coefficient, the shape of the detail coefficient is shown in Fig. 4. After all messages are inserted into the detail coefficients (cD), the next step is to return the signal to the time domain so that it can be represented as an audio signal.

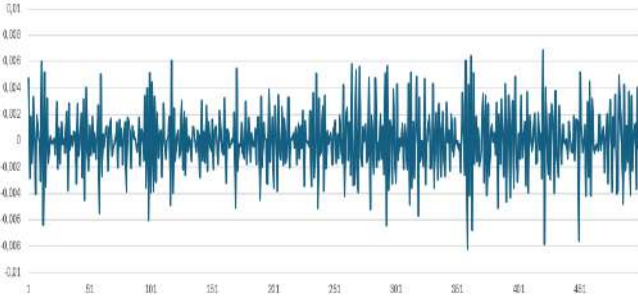


Fig. 4. Detail Coefficient that has embedded secret message

C. Message Extraction Process

The audio signal in .wav format is read and represented as digital amplitude values. Next, a DWT is applied to obtain cA and cD. The message extraction process is carried out by analyzing the embedded audio signal in the detail coefficient domain (cD). Each cD frame is processed with the Hann Window, and its cepstrum is computed to obtain a representation in the quefrency domain. The presence of echo is indicated by a peak at a certain delay. If the peak appears around the X_0 position, it represents bit '0', while if it appears around X_1 , it represents bit '1'. Thus, the secret message can be reconstructed from the audio signal. An example of the cepstrum calculation results is shown in Fig. 5. In the test results, the main peak is detected near the X-axis position ≈ 261 . This position indicates the presence of an echo component with a specific delay embedded in the signal. This position is closer to X_1 than X_0 , so it can be concluded that the frame represents bit '1'.

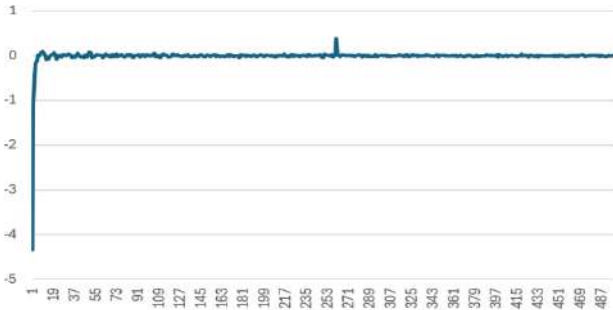


Fig. 5. Cepstrum Calculation Results

D. Performance Evaluation and Robustness Testing

The performance evaluation was conducted to assess imperceptibility, extraction accuracy, and robustness of the proposed method. Imperceptibility was measured using SNR, and extraction accuracy using BER. The evaluation includes clean-condition testing, alpha parameter analysis, baseline comparison, Gaussian noise attack, and MP3 compression attack.

Testing was conducted to evaluate the performance of the developed steganography method. Six WAV audio files (see Table II) were used as cover audio. Two secret messages were used for embedding, as shown in Table III.

TABLE III. MESSAGE FOR EMBEDDING

#	Message
Message 1	The quick brown fox jumps over the lazy dog
Message 2	String with ASCII codes ranging from 0 to as many as the cover audio can hold.

Each message will be embedded in each cover audio file. After embedding, the Signal-to-Noise Ratio (SNR) value is calculated to measure the imperceptibility of the stego audio to the original audio. The results of this calculation are used to analyze the extent to which the secret message embedding affects audio quality and to ensure that the message can be embedded without significantly reducing it. The test findings are displayed in Table IV.

TABLE IV. SNR, BER VALUE OF STEGO AUDIO WITH PARAMETERS $X_0=60$, $X_1=260$, ALPHA=0.7

Experiment	Cover Audio	Message 1		Message 2	
		SNR (dB)	BER	SNR (dB)	BER
1.	Cover Audio-1	47.11	0.00	39.69	0.00
2.	Cover Audio-2	32.35	0.00	28.36	0.00
3.	Cover Audio-3	24.24	0.00	19.92	0.00
4.	Cover Audio-4	20.96	0.00	15.96	0.00
5.	Cover Audio-5	49.92	0.00	45.45	0.00
6.	Cover Audio-6	48.73	0.00	42.65	0.00

The SNR values generated by embedding the two messages ranged from 15.96 dB to 49.92 dB. Although the lowest SNR was 15.96 dB, subjective testing using the Mean Opinion Score (MOS) method with 10 respondents yielded an average of 4.8, indicating that the audio quality was still in the good category. This condition indicates that the message embedding did not significantly interfere with the listener's perception. In addition, the insertion of the detail coefficient (cD) and structured echo characteristics caused the signal changes to be perceived as non-disturbing noise.

Testing by embedding message 1 into cover audio-4 and cover audio-5 produces significantly different SNR values. The SNR value for audio cover-4 is 20.96 dB, while for audio cover-5 it is 49.92 dB. This can occur because the SNR is greatly influenced by the characteristics of the audio cover, not just the algorithm. Signals with loud audio, or in other words, large amplitude, will produce high SNR values. Conversely, low-amplitude audio signals will yield low SNR values. Therefore, the SNR value is also influenced by the strength of the signal from the audio cover.

In addition to low amplitudes, small SNR values can also be affected by the α (alpha) parameter. The larger the α value, the stronger the echo produced, thereby increasing the difference between the original and stego signals. This results in a lower SNR. Conversely, small α values tend to produce higher SNR values, as the changes occurring in the signal are relatively smaller. However, increasing the α value will

strengthen the echo, thereby increasing the success rate of secret message retrieval.

At $\alpha = 0.7$, all messages can be extracted correctly (BER = 0). This indicates that the echo strength is sufficient to form a clear delay pattern. However, when α is lowered to 0.6, extraction errors begin to occur (BER > 0). This is caused by the reduced echo amplitude, which makes it less dominant over the original signal, especially in the low-signal region. This indicates that decreasing α improves audio quality but reduces the system's robustness. The SNR and BER values with $\alpha = 0.6$ are shown in Table V.

TABLE V. SNR, BER VALUE OF STEGO AUDIO WITH PARAMETERS $X_0=60$, $X_1=260$, ALPHA=0.6

Experiment	Cover Audio	Message 1		Message 2	
		SNR (dB)	BER	SNR (dB)	BER
1.	Cover Audio-1	49.34	0.00	39.7	0.00
2.	Cover Audio-2	34.60	0.00	28.4	0.01
3.	Cover Audio-3	26.48	0.00	19.9	0.00
4.	Cover Audio-4	23.22	0.00	16.0	0.00
5.	Cover Audio-5	52.03	0.02	45.5	0.01
6.	Cover Audio-6	50.91	0.00	42.7	0.00

In Cover Audio-5, the BER value for message 1 is 0.02. This indicates a failure to retrieve the embedded message. TABLE VI shows the differences between message 1 and the extracted message. The failure occurs when the space character between the words "the" and "lazy" is retrieved. The resulting character is a null character.

TABLE VI. DIFFERENCE BETWEEN ORIGINAL MESSAGE AND EXTRACTED MESSAGE WITH ALPHA=0.6

Original Message	The quick brown fox jumps over the lazy dog
Extracted Message	The quick brown fox jumps over the lazy dog

To further validate the effectiveness of the proposed method, a comparative evaluation was conducted using conventional audio steganography techniques under the same experimental conditions. The comparison includes Least Significant Bit (LSB), standalone Echo Hiding, DWT-based LSB steganography, and the proposed DWT + Echo Hiding method using the same cover audio files, secret message, and evaluation metrics. The comparison results in terms of Signal-to-Noise Ratio (SNR) and Bit Error Rate (BER) are presented in TABLE VII.

TABLE VII. PERFORMANCE COMPARISON OF AUDIO STEGANOGRAPHY METHODS

Method	Average SNR (dB)	Average BER
LSB	107.6763	0.000000
Echo Hiding	6.8433	0.000706
DWT + LSB	65.2592	0.102843
DWT + Echo Hiding	34.6117	0.000000

As shown in TABLE VII, LSB achieved the highest average SNR and zero BER under clean extraction conditions. However, the proposed DWT + Echo Hiding method also achieved zero BER while maintaining a moderate average SNR. This indicates that the proposed method does not outperform LSB in terms of imperceptibility under non-attack conditions. Nevertheless, the proposed method provides a different balance by combining echo-based embedding with wavelet-domain coefficient modification. Therefore, the contribution of this study should be interpreted as a balanced implementation of Echo Hiding in DWT detail coefficients rather than a universal superiority over all conventional methods.

To further evaluate the robustness of the proposed method, the stego audio generated using the selected embedding parameters ($X_0 = 60$, $X_1 = 260$, and $\alpha = 0.7$) was subjected to additive white Gaussian noise (AWGN). Different values of the noise standard deviation (σ) were used to simulate channel distortion prior to message extraction. The resulting Bit Error Rate (BER) values under each noise condition are presented in TABLE VIII.

TABLE VIII. BER OF THE PROPOSED METHOD UNDER GAUSSIAN NOISE ATTACK

Stego Audio	Standard Deviation (σ)				
	0.00001	0.00005	0.00010	0.00050	0.00100
Stego Audio-1	0.00000	0.01483	0.03072	0.17373	0.26801
Stego Audio-2	0.00318	0.00318	0.01271	0.03920	0.04661
Stego Audio-3	0.00000	0.00000	0.00000	0.01271	0.02966
Stego Audio-4	0.00530	0.00847	0.00845	0.01271	0.02436
Stego Audio-5	0.02966	0.12923	0.25212	0.40466	0.40466
Stego Audio-6	0.00212	0.01377	0.03708	0.24788	0.32309

As shown in TABLE VIII, the BER generally increased as the noise standard deviation (σ) increased, indicating that higher levels of additive white Gaussian noise (AWGN) progressively degraded the embedded echo pattern and reduced the accuracy of cepstrum-based delay detection. At low noise levels ($\sigma \leq 0.0001$), most cover audio files still achieved low BER values, indicating that the proposed method is tolerant of mild channel distortion. However, as the noise intensity increased to $\sigma = 0.0005$ and $\sigma = 0.001$, the BER increased noticeably, indicating that excessive noise interferes with the embedded echo signal and reduces the reliability of message extraction.

To evaluate the robustness of the proposed method against lossy compression, all stego audio files generated from six cover audio files and two secret messages were compressed to MP3 at 128 kbps, then converted back to WAV before extraction. This test was conducted to examine whether the embedded echo pattern in the DWT detail coefficients could still be detected after the compression-decompression process. The BER values obtained under the MP3 compression attack are presented in TABLE IX.

TABLE IX. BER OF THE PROPOSED METHOD UNDER MP3 COMPRESSION ATTACK

Cover Audio	BER
-------------	-----

	Message 1	Message 2
Cover Audio-1	0.0731	0.0678
Cover Audio-2	0.0381	0.0350
Cover Audio-3	0.0350	0.0286
Cover Audio-4	0.0403	0.0297
Cover Audio-5	0.1335	0.1186
Cover Audio-6	0.0169	0.0064
Average	0.0561	0.0477

As shown in TABLE IX, MP3 compression introduced extraction errors in all tested stego audio files. The average BER values were 0.0561 for Message 1 and 0.0477 for Message 2. The lowest BER was observed with Cover Audio-6, whereas Cover Audio-5 produced the highest BER for both messages. These results indicate that the proposed method has partial tolerance against MP3 compression, but it cannot guarantee error-free extraction after lossy compression. The variation in BER values also shows that compression robustness is influenced by the characteristics of the cover audio.

IV. CONCLUSION

This research proposed an audio steganography method that integrates Echo Hiding into the detail coefficients of the Discrete Wavelet Transform (DWT). The proposed method was evaluated using six WAV audio files and two secret messages. The experimental results show that the selected parameters, $X_0 = 60$, $X_1 = 260$, and $\alpha = 0.7$, achieved error-free message extraction under clean conditions, with BER = 0 for all tested audio files and messages. The obtained SNR values ranged from 15.96 dB to 49.92 dB, while the subjective evaluation produced an average MOS of 4.8, indicating that the stego audio quality remained perceptually acceptable.

The results confirm that parameter selection significantly affects the balance between imperceptibility and extraction accuracy. Reducing α to 0.6 increased the SNR values in several cases, but extraction errors began to occur, showing that weaker echo embedding may reduce distortion while also reducing extraction reliability. The robustness evaluation under AWGN and MP3 compression shows that the proposed method has partial tolerance against selected signal-processing attacks, but it cannot guarantee error-free extraction under distorted conditions. In particular, MP3 compression at 128 kbps produced average BER values of 0.0561 for Message 1 and 0.0477 for Message 2.

Overall, the proposed method provides a practical balance between audio quality and extraction accuracy under the evaluated conditions. However, the robustness evaluation is still limited to Gaussian noise and MP3 compression attacks. Future work should include filtering and resampling attacks, larger, more diverse audio datasets, statistical analysis, and comparative benchmarking under identical attack conditions. Further improvements may also consider adaptive synchronization, stronger echo coding, and error-

correction mechanisms to improve message recovery under stronger signal distortions.

REFERENCES

- [1] O. Evsutin, A. Melman, and R. Meshcheryakov, 'Digital steganography and watermarking for digital images: A review of current research directions', *IEEE Access*, vol. 8, pp. 166589–166611, 2020, doi: 10.1109/ACCESS.2020.3022779.
- [2] M. B. Yel and M. K. M. Nasution, 'KEAMANAN INFORMASI DATA PRIBADI PADA MEDIA SOSIAL', *Jurnal Informatika Kaputama*, vol. 6, no. 1, pp. 92–101, Jan. 2022.
- [3] F. Hazzaa, A. M. Shabut, N. H. M. Ali, and M. Cirstea, 'Security Scheme Enhancement for Voice over Wireless Networks', *Journal of Information Security and Applications*, vol. 58, p. 102798, May 2021, doi: 10.1016/j.jisa.2021.102798.
- [4] S. S. Radke and D. S. Mishra, 'Review of Image Security approaches: Concepts, Issues, Challenges and Applications', *Journal of University of Shanghai for Science and Technology*, vol. 23, pp. 1047–1054, Jun. 2021.
- [5] Musa. M. Yahaya and A. Ajibola, 'Cryptosystem for Secure Data Transmission using Advance Encryption Standard (AES) and Steganography', *International Journal of Scientific Research in Computer Science, Engineering and Information Technology*, pp. 317–322, Dec. 2019, doi: 10.32628/cseit195659.
- [6] A. K. Agrahari, M. Sheth, and N. Praveen, 'Comprehensive Survey on Image Stegnography Using LSB With AES In-depth study of various requirements in the process of embedding, encrypting and extracting of text data', 2018. [Online]. Available: <http://www.ripublication.com>
- [7] R. K. Dewi and R. Munir, 'PERBANDINGAN BERBAGAI METODE STEGANOGRAFI PADA CITRA DIGITAL', *Jurnal Informatika Polinema*, vol. 9, no. 3, pp. 289–230, May 2023, doi: 10.33795/jip.v9i3.1336.
- [8] F. Yanti and K. Budayawan, 'Jurnal Vocational Teknik Elektronika dan Informatika', *Jurnal Vocational Teknik Elektronika dan Informatika*, vol. 11, no. 1, Mar. 2023, [Online]. Available: <http://ejournal.unp.ac.id/index.php/voteknika/index>
- [9] A. Dwi Hendrata and A. Prihanto, 'Analisis Kualitas Suara Stego Audio Penyisipan Informasi Tersembunyi dengan Metode Least Significant Bit', *Journal of Informatics and Computer Science*, vol. 02, 2021, [Online]. Available: <https://www2.cs.uic.edu/~i101/SoundFiles/>.
- [10] N. Kaur and S. Behal, 'Audio Steganography Techniques-A Survey', 2014. [Online]. Available: www.ijera.com
- [11] S. A. Br Sibarani, A. Munthe, R. Purba, and A. A. Lubis, 'Pengamanan Citra Digital Menggunakan Kriptografi DNA dan Modified LSB', *Jurnal Teknologi Informasi dan Ilmu Komputer*, vol. 11, no. 6, pp. 1233–1242, Dec. 2024, doi: 10.25126/jtiik.2024117666.
- [12] L. F. Adhimah, I. Nurhafiyah, A. A. Muntahar, F. Kristiaji, and D. Mustofa, 'IMPLEMENTASI APLIKASI STEGANOGRAFI BERBASIS WEB MENGGUNAKAN ALGORITMA LSB DAN

- [13] S. Rahman *et al.*, 'A novel and efficient digital image steganography technique using least significant bit substitution', *Nature Research*, Dec. 2025. doi: 10.1038/s41598-024-83147-3.
- [14] C. E. Pangaribuan, O. K. Surbakti, F. Sihotang, Y. Tamba, Y. D. Ndruru, and M. Hasugian, 'Analysis Of The Utilization Of Echo Hiding Algorithm In Hiding Secret Messages', *Journal Majelis Paspama*, vol. 2, no. 2, pp. 102–107, 2024.
- [15] S. Mishra, V. K. Yadav, M. C. Trivedi, and T. Shrimali, 'Audio steganography techniques: A survey', in *Advances in Intelligent Systems and Computing*, Springer Verlag, 2018, pp. 581–589. doi: 10.1007/978-981-10-3773-3_56.
- [16] S. Wen, H. Tan, J. Li, and T. Hu, 'A robust audio watermarking method based on dual-encoder U-Net and short-time Fourier transform', *Cyber Security and Applications*, vol. 4, Jan. 2026, doi: 10.1016/j.csa.2026.100129.
- [17] I. Putu *et al.*, 'Perancangan Sistem Steganografi Berbasis Transformasi Wavelet Diskrit Terintegrasi Algoritma Rijndael dan QR-Code', *JNATIA*, vol. 2, no. 4, 2024.
- [18] R. F. Damanik, R. Ardhia Pramesthi, G. Budiman, and S. Saidah, *Audio Watermarking Berbasiskan DWT-DCT-SVD-CPT Menggunakan QIM dan SS Audio Watermarking Based on DWT-DCT-SVD-CPT Using QIM and SS*. 2019.
- [19] G. A. Simanungkalit, D. S. Tarigan, and D. R. Simangunsong, 'Discrete Wavelet Transform (DWT) Based Steganography Implementation', 2023. [Online]. Available: <https://jurnal.seaninstitute.or.id/index.php/jutip13>
- [20] C. Umam and M. Muslih, 'Enkripsi Data Teks Dengan AES dan Steganografi DWT', *InComTech: Jurnal Telekomunikasi dan Komputer*, vol. 13, no. 1, p. 28, Apr. 2023, doi: 10.22441/incomtech.v13i1.15059.
- [21] A. Syahdilan and M. A. Prawira, 'IMPLEMENTASI ALGORITMA DISCRETE WAVELET TRANSFORM UNTUK MENAMBAHKAN INVISIBLE WATERMARKING PADA CITRA DIGITAL', 2024.
- [22] H. Rafiee and M. Fakhredanesh, 'Presenting a Method for Improving Echo Hiding', *Journal of Computer and Knowledge Engineering*, vol. 2, no. 1, 2019, doi: 10.22067/cke.v2i1.74388.
- [23] S. Lahiri, 'Audio Steganography using Echo Hiding in Wavelet Domain with Pseudorandom Sequence', *Int. J. Comput. Appl.*, vol. 140, no. 2, pp. 16–22, Apr. 2016, doi: 10.5120/ijca2016909223.
- [24] N. Faradiba, 'Frekuensi dan Intensitas Suara yang Bisa Didengar Manusia', <https://www.kompas.com/sains/read/2022/06/18/123100923/frekuensi-dan-intensitas-suara-yang-bisa-didengar-manusia>, Jun. 18, 2022.
- [25] M. Hamdani and G. N. Samosir, 'IMPLEMENTASI STEGANOGRAFI UNTUK KEAMANAN PENGIRIMAN CITRA DIGITAL MENGGUNAKAN METODA DCT (DISCRETE COSINE TRANSFORM)', *Jurnal Penelitian dan Pengkajian Elektro*, vol. XX, no. 2, p. 42, Apr. 2018.

Lampiran 4. HKI

 REPUBLIK INDONESIA KEMENTERIAN HUKUM	
SURAT PENCATATAN CIPTAAN	
Dalam rangka perlindungan ciptaan di bidang ilmu pengetahuan, seni dan sastra berdasarkan Undang-Undang Nomor 28 Tahun 2014 tentang Hak Cipta, dengan ini menerangkan:	
Nomor dan tanggal permohonan	: EC002026130446, 31 Juli 2026
Pencipta	
Nama	: Mardi Hardjianto, Jeffrey Bram Pattipeilohy dkk
Alamat	: Jl. Meneng II/31, RT.012 RW.007, Grogol Petamburan, Kota Adm. Jakarta Barat, DKI Jakarta, 11470
Kewarganegaraan	: Indonesia
Pemegang Hak Cipta	
Nama	: Direktorat Riset dan Pengabdian kepada Masyarakat
Alamat	: Jl. Ciledug Raya, RT.10/RW.2, Pesanggrahan, Kota Adm. Jakarta Selatan, DKI Jakarta, 12260
Kewarganegaraan	: Indonesia
Jenis Ciptaan	: Program Komputer
Judul Ciptaan	: Steganografi Audio Menggunakan DWT dan Echo Hiding
Tanggal dan tempat diumumkan untuk pertama kali di wilayah Indonesia atau di luar wilayah Indonesia	: 1 Juli 2026, di DKI Jakarta
Jangka waktu perlindungan	: Berlaku selama 50 (lima puluh) tahun sejak Ciptaan tersebut pertama kali dilakukan Pengumuman.
Nomor Pencatatan	: 001389057
adalah benar berdasarkan keterangan yang diberikan oleh Pemohon. Surat Pencatatan Hak Cipta atau produk Hak terkait ini sesuai dengan Pasal 72 Undang-Undang Nomor 28 Tahun 2014 tentang Hak Cipta.	
	a.n. MENTERI HUKUM DIREKTUR JENDERAL KEKAYAAN INTELEKTUAL u.b Direktur Hak Cipta dan Desain Industri  Agung Damarsasongko,SH.,MH. NIP. 196912261994031001
	Disclaimer: 1. Dalam hal pemohon memberikan keterangan tidak sesuai dengan surat pernyataan, Menteri berwenang untuk mencabut surat pencatatan permohonan. 2. Surat Pencatatan ini telah disegel secara elektronik menggunakan segel elektronik yang diterbitkan oleh Balai Besar Sertifikasi Elektronik, Badan Siber dan Sandi Negara. 3. Surat Pencatatan ini dapat dibuktikan keasliannya dengan memindai kode QR pada dokumen ini dan informasi akan ditampilkan dalam browser.



LAMPIRAN PENCIPTA

No	Nama	Alamat
1	Mardi Hardjianto	Jl. Menteng II/ 31, RT 012 RW 007 Grogol Petamburan, Kota Adm. Jakarta Barat
2	Jeffrey Bram Pattipeilohy	Komplek Pohri Ulujam Blok H-1 Pesanggrahan, Kota Adm. Jakarta Selatan
3	Dewi Kusumaningsih	Jl. Bungur RT.002/RW.011 No. 55B Pesanggrahan, Kota Adm. Jakarta Selatan
4	Blasius Dala Nai	Jl. Palmerah barat no. 21 RT 004/005, Grogol Utara dan kode pos Jakarta 12210 Grogol Petamburan, Kota Adm. Jakarta Barat



Disclaimer:

1. Dalam hal pemohon memberikan keterangan tidak sesuai dengan surat pernyataan, Menteri berwenang untuk mencabut surat pencatatan permohonan.
2. Surat Pencatatan ini telah disegel secara elektronik menggunakan segel elektronik yang diterbitkan oleh Balai Besar Sertifikasi Elektronik, Badan Siber dan Sandi Negara.
3. Surat Pencatatan ini dapat dibuktikan keasliannya dengan memindai kode QR pada dokumen ini dan informasi akan ditampilkan dalam browser.