

Audio Watermarking with Patchwork and DCT

Yugendra D. Chincholkar

Department of Electronics and Telecommunication, Sinhgad College of Engineering, Pune, India
ydchincholkar.scoe@sinhgad.edu

Sanjay R. Ganorkar

Department of Electronics and Telecommunication, Sinhgad College of Engineering, Pune, India
srganorkar.scoe@sinhgad.edu

Article Info

Volume 83

Page Number: 2159 - 2167

Publication Issue:

March - April 2020

Abstract

Digital audio watermarking plays a crucial role in preserving multimedia content from illegitimate copying and reproduction. Audio watermarking has a relatively good perception quality level for the patchwork process. The concerns between protection, robustness, and imperceptibility are contemporary research areas, which are still significant issues. Proposed research explores, a patchwork approach for signal processing attacks, using discrete cosine transforms (DCT) within audio watermarking slant. At first, the watermarking audio clip is fragmented up into alike portion and their sub-portions, followed by its coefficient calculation. After removing the coefficient associated with high frequency, the residual coefficients are utilized for the formation of equal length of frame pairs. Using specific embedding criteria and secure information, key watermarks are inserted into audio. The adjustments are made pertaining, the selection conditions chosen through the embedding procedure, so that, the recognition of watermarked pair of frames, is made at the extraction step. To remove the watermark from watermarked audio, the private data key is used. The explored audio watermarking technique is applied and evaluated for security, robustness, security, imperceptibility, and data payload under signal processing attacks.

Article History

Article Received: 24 July 2019

Revised: 12 September 2019

Accepted: 15 February 2020

Publication: 18 March 2020

Keywords – Discrete cosine transform Patchwork, Robustness, Imperceptibility, Security.

I. INTRODUCTION

Gradually, Internet technologies are increasingly being used for the suffering and downloading of online multimedia data; this is due to developments in electronic gadgets with internet technology. The storage, broadcasting, manipulation, and circulation of electronic implements in multimedia data are, therefore, possible without any deterioration in their performance. To avoid unethical access, which leads to a huge urge for secured data? Digital watermarking is a vital innovation for securing rights and authenticating honor in open network environments [1]. Strictly, the purpose of digital watermarking is to retain undisclosed information about watermarking, in the authentic media data, by keeping its frequent use. If it is appropriate to remove watermarks data [2]– [3], holders may claim

Copyright. Generally, electronic watermarking is categorized as audio, video, and image watermarking [2][3][5] based on the appropriate areas. This paper's primary focus is on the phenomenon of audio watermarking. The perception that is perceived by human echoes is to a greater extent sensitive than the sensor's alternative perceptions, such as vision [2][3][5]. It is possible to conceal extra information in perceptible signals devoid of reducing the superiority of mass medium data compared to other multimedia information or data. Inaudibility, robustness, and security are three significant characteristics of an effectively competent and practical audio watermarking scheme. The practically inaudible knowledge of the implanted watermark is the imperceptibility. Robustness refers to the capability of gathering watermark statistics from the watermarked signal. Criteria for robustness and imperceptibility are

discordant but vital to be met. Security means to prevent unauthorized people from removing watermarks apart from the secret key, during the watermarking process by employing secured Key. Besides these features, the watermarking process also has other advantages of low computational burden with improvements. An effective watermarking scheme is especially meant for time-consuming applications and a custom watermarking scheme integrates different applications. It is possible to excerpt the embedded watermarks at the decoding point instead of the host's audio data [1]. It's attractive, too. Such an audio watermarking type is referred to as blind audio watermarking. In current years, several watermarking approaches have been offered for audio signals such as Least significant it, Vector regression support [3][5], Echo hiding [2][5], Spread Spectrum, Patchworks [1][5]. Execution of audio watermarking with the Patchwork-based method is highly robust and has the excessive ability to sustain, against attacks and offer a high degree of imperceptibility, security and robustness can also be achieved.

Bender et al. [5] predicted the patchwork method to advanced signal processing attacks as an emerging method for watermarking with strong robustness. Laurence Boney et al. [6] presented the watermarking method for audio for the first time in 1996. Arnold then extended this technique to audio signals [7]. Kwon Yeo et al. [8] proposed an improved patchwork process, using domain conversions but failed to reach a high bit speed. H. Kang et al. [9] suggested an unsighted process for watermarking, based on a full index blend, to outperform in a psychologically adjustable based patchwork methodology. N. K. Kalantari et al. [10] has suggested a flexible multiplicative patchwork approach to audio watermarking. However, refuses to target forward signal processing. The Chi-Man Pun et al. audio watermarking algorithm [11] proposed and implemented in two phases using an adaptive patchwork approach with Neural Network, it could not sustain with advanced attacks. A system

based on patchwork encryption and decoding, with improved performance, was explored by I. Natgunanathan et al. [12]. A mathematical analysis was discussed with the theoretically determined values for the assessment of practical audio signal measurements in [13]. Peng Hu et al. [15] suggested an improved audio watermarking algorithm utilizing the Constant Q Transform (CQT) field. Audio performance perceptual analysis (PEAQ) and Bit Error Rate is managed as reliability variables in the scale of 0 – 1.2 percent to exhibit the cogency of the algorithm. In addition to the findings of similarities with the high-fidelity signal, I. Natgunanathan et al. [14] recommended for a stereo signal in the frequency domain, a novel and most upgraded phase of a digital audio watermarking process. It is not a robust solution to sophisticated attacks like de-sync, pitch, or change in time. However, it is seen that much of the research is discussed, under certain attacks and yet no researcher has considered all attacks together as they need to meet audio watermarking characteristics such as robustness, security, data payload, and imperceptibility.

We have proposed a watermarking framework based on audio signal patchwork. Here, the host audio portion is subdivided into two sub-sections and then DCT is determined for individual sub-section coefficient. Further, the coefficients of the high-frequency variable DCT are eliminated, then the remainder is split into numerous slabs. Certain measures are used to select pairs of DCT frames suitable to implant watermarks. The pairs of the DCT framework are separated into many frames. Further, a watermark is applied on selected pairs after updating the relevant DCT coefficients as a synchronization code with the help of a pseudo-noise (PN) sequence. Further, certain selection measures are used in the embedding stage, an integration algorithm is proposed to identify watermark blocks in the watermark signal. With the PN sequence and the secret key, watermarks can be extracted after the system pairs with watermarks have been found. The patchwork method proposed

is fairly improved to existing watermarking techniques, as the watermarked block at the decoding end does not require extra data and are exceptionally resilient to attacks. The simulation results show their efficiency.

The paper's respite was organized correspondingly. Section I I presents the state-of-the-art encryption and decoding system based on patchwork. In Section III, the simulation results to validate the success of the development. The interpretation is described in Section IV.

II. EMBEDDING AND DECODING PROCESS

A. Watermark Inserting Process

First, the selected audio clip is filtered and separated into an identical number of sub-portions with the fixed-length M to reach a range of sub-portion Where M is the even number and real. Let $\mathcal{R}(n)$ be the 60-second audio clip. Let $\mathcal{R}(n)'$ be the sub-portion of the audio clip. Each sub-portion of the audio clip is further separated into sub-sections of equal length R where R is even number and real and. The following equation (1) is then used to determine the DCT coefficients of each subsection.

$$\mathcal{R}(\zeta) = \frac{1}{\sqrt{R}} \sum_{n=0}^{L-1} \mathcal{R}(n)' \cos \left\{ \frac{\pi(2n+1)\zeta}{2R} \right\} \quad (1)$$

Where $\zeta = 0, 1, \dots, R-1$

DCT is a digital conversion which converts real information facts into their actual frequency range. It shirks the problem of duplication because DCT can reduce the energy of the spatial series to minimum frequency coefficients. Later then We realize that the coefficients associated with high frequency are vulnerable to signal processing attacks and hereafter not being included during the implanting phase of the watermark. For the implanting of watermarks, only selected frequency (i.e. low and mid-frequency) coefficients of each sub-section are regarded. Let $\mathcal{R}_s(\zeta)$ be the selected coefficient of the frequency range from $\mathcal{R}(\zeta)$. After elimination, each sub-sections of low and mid-frequency coefficients are

considered to create a frame composed of an identical number of coefficients. Let $\mathcal{R}_s(\zeta)$ further, divide the equal number of P_{LS} length frames $\mathcal{R}_{s,K}(\zeta)$, $K = 1, 2, \dots, P_{LS}$ frames.

$$\mathcal{R}(\zeta) = [\mathcal{R}_{s,1}(\zeta), \mathcal{R}_{s,2}(\zeta), \dots, \mathcal{R}_{s,P_{LS}}(\zeta)] \quad (2)$$

However, the individual frame is split into subframes of the same length $\mathcal{R}_{s,K}(k)$ as follows:

$$\mathcal{R}_{s,K,1}(\zeta) = [\mathcal{R}_{s,K}(0), \mathcal{R}_{s,K}(1), \dots, \mathcal{R}_{s,K}(P_{LS}/2 - 1)] \quad (3)$$

$$\mathcal{R}_{s,K,2}(\zeta) = [\mathcal{R}_{s,K}(P_{LS}/2), \mathcal{R}_{s,K}(P_{LS}/2 + 1), \dots, \mathcal{R}_{s,K}(P_{LS} - 1)] \quad (4)$$

Therefore, each frame consists of

$$\mathcal{R}_{s,K}(\zeta) = [\mathcal{R}_{s,K,1}(\zeta), \mathcal{R}_{s,K,2}(\zeta)] \quad (5)$$

When creating a set of frames, for security and synchronization purposes, a PN sequence $\ell(n) = \{\ell(1), \ell(2), \ell(3), \dots, \ell(M)\}$ of the frame size length is generated. This generated sequence scrambled with frames coefficients and formed a synchronization code frame. To offer high-level security, which is important parameters of watermarking, private key $\mathcal{L}(n)$ of the one-dimensional sequence is used, and it is scrambled with watermark bit $h(0)$. Let $\omega(n)$ be the watermark data produced by the scrambling process

$$\omega(n) = \mathcal{L}(n) \oplus h(0) \quad (6)$$

Where $\oplus$ represents mod 2 operation. They get mixed with DCT frames once the scrambled data is generated.

This scrambled watermark data embedded into the synchronization code frame. Several numbers of frame pairs cannot be used for the implantation of watermarks in the framework pairs. Inserting the watermark in these frame pairs may decrease the

watermarked signal perceptual quality to an undesirable level. To select the proper synchronization code frame pairs $|\mathcal{R}_{s,K,1}(\zeta)|$ and $|\mathcal{R}_{s,K,2}(\zeta)|$, We define and calculate the mean value of such frames by utilizing equation (7).

$$m_{K,1} = E(|\mathcal{R}_{s,K,1}(\zeta)|)$$

$$m_{K,2} = E(|\mathcal{R}_{s,K,2}(\zeta)|)$$

(7)

$$m_K = E(|\mathcal{R}_{s,K}(\zeta)|)$$

$$n_K = \frac{m_{K,1} + m_{K,2}}{2}$$

The K^{th} DCT sub-section is referred to as an un-silent sub-section when the absolute value of the

DCT sub-section implies $m_K > \partial$. We use all non-silent DCT sub-section to add watermarks to increase the effectiveness of implanting and ensure commendable perceptual accuracy with constraints.

$$m_K > \partial \quad \text{and} \quad n_K \geq g \quad (8)$$

Where $K = 1, 2, \dots, P_{LS}$ and ∂ and g is an arbitrarily selected threshold value for robustness and imperceptibility control.

As soon as the above condition is mention in equation (8) is satisfied, the scrambled watermark $\omega(n)$ bits will only be inserted into the synchronization code frame using specific rules. The frame coefficients are altered in terms of mean value after adding watermark bits.

$m'_{K,1}$ and $m'_{K,2}$

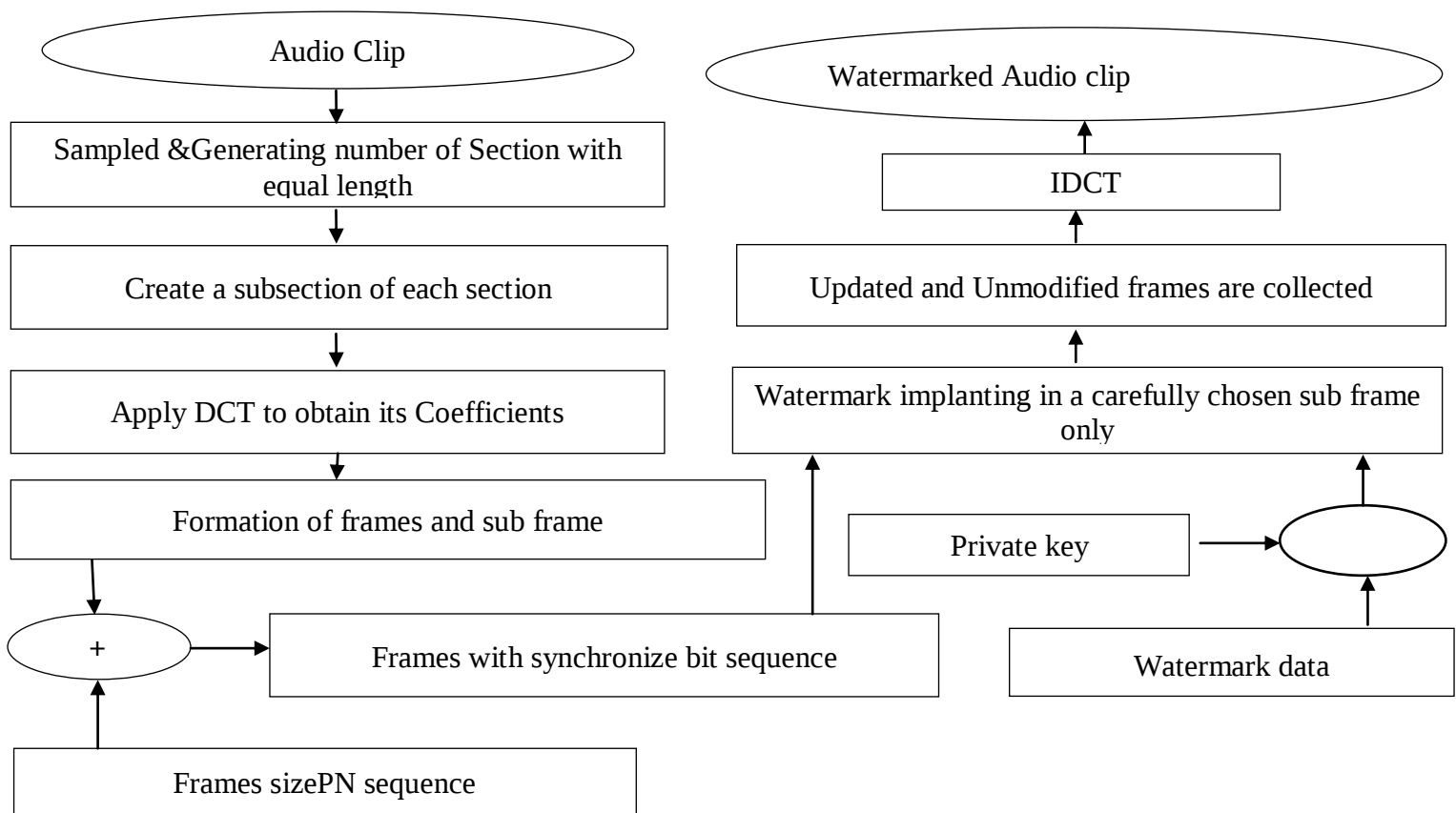


Fig. 1. Watermark implanting method in the audio clip

For each K^{th} the coefficients of the scrambled subsegment are modified using the following equation:

$$\mathcal{R}'_{s,K,1}(\zeta) = \mathcal{R}_{s,K,1}(\zeta) \times \frac{m'_{K,1}}{m_{K,1}} \quad \text{and}$$

$$\begin{aligned} \mathcal{R}'_{s,K,2}(\zeta) \\ \times \frac{m'_{K,2}}{m_{K,2}} \end{aligned} = \mathcal{R}_{s,K,2}(\zeta) \quad (10)$$

Then it is possible to obtain all counterparts marked with the watermark:

$$\begin{aligned} \mathcal{R}'_{s,K,1}(\zeta) = [\mathcal{R}'_{s,K,2}(\zeta), \mathcal{R}'_{s,K,2}(\zeta)] \quad , K \\ = 1, 2, \dots, P_{LS} \end{aligned} \quad (11)$$

and

$$\begin{aligned} \mathcal{R}'_s(\zeta) \\ = [\mathcal{R}'_{s,1}(\zeta), \mathcal{R}'_{s,2}(\zeta), \dots, \mathcal{R}'_{s,P_{LS}}(\zeta)] \end{aligned}$$

In an identical way, we can gather the same watermark bit from the same host audio sub-section in all selected DCT frame pairs. After that, we are forming the watermarked audio clip employing the Inverse Discrete Cosine Transform (IDCT).

B. Watermark Extraction Process

Once the watermark audio clip has been received, the subsequent step is to divide the watermarked audio clip into several sections. The division of the audio clip is done similarly used in the implanting process. We implement DCT to the audio clip sub-sections to obtain their coefficients. Let $\mathcal{R}(\zeta)'$ be the of the sub-section DCT coefficients.

Let $\mathcal{R}_s(\zeta)'$ is further separated into the equal number of frames of equal P_{LS} length $\mathcal{R}_{s,K}(\zeta)$, $K = 1, 2, \dots, P_{LS}$

$$\begin{aligned} \mathcal{R}_s(\zeta)' \\ = [\mathcal{R}_{s,1}(\zeta)', \mathcal{R}_{s,2}(\zeta)', \dots, \mathcal{R}_{s,P_{LS}}(\zeta)'] \end{aligned}$$

Each frame is getting additional divided into subframes of equal length $\mathcal{R}_{s,i}(k)'$ as follows:

$$\begin{aligned} \mathcal{R}_{s,K,1}(\zeta)' \\ = [\mathcal{R}_{s,K}(0)', \mathcal{R}_{s,K}(1)', \dots, \mathcal{R}_{s,K} \left(\frac{P_{LS}}{2} - 1 \right)'] \end{aligned} \quad (14)$$

$$\begin{aligned} \mathcal{R}_{s,K,2}(\zeta)' \\ = [\mathcal{R}_{s,K} \left(\frac{P_{LS}}{2} \right)', \mathcal{R}_{s,K} \left(\frac{P_{LS}}{2} + 1 \right)', \dots, \mathcal{R}_{s,K}(P_{LS} - 1)'] \end{aligned} \quad (15)$$

(12)
Thus, each frame contains

$$\begin{aligned} \mathcal{R}_{s,K}(\zeta)' \\ = [\mathcal{R}_{s,K,1}(\zeta)', \mathcal{R}_{s,K,2}(\zeta)'] \end{aligned} \quad (16)$$

After fixing a group of frames, A sequence of PNof a frame size length is produced for protection and synchronization. When attempting to remove a watermark from a pair of frames, it is necessary to find out the watermark contains a pair of frames. It is therefore essential to find the means value of the pairs of frames $|\mathcal{R}_{s,K,1}(\zeta)'|$ and $|\mathcal{R}_{s,K,2}(\zeta)'|$.

$$\begin{aligned} m_{K,1}' &= E(|\mathcal{R}_{s,K,1}(\zeta)'|) \\ m_{K,2}' &= E(|\mathcal{R}_{s,K,2}(\zeta)'|) \\ (17) \quad m_K' &= E(|\mathcal{R}_{s,K}(\zeta)'|) \end{aligned}$$

$$n_K' = \frac{m_{K,1}' + m_{K,2}'}{2}$$

If $m_K' > \partial$ and $n_K' \geq \phi$ is fulfilled, which shows that the particular pair of frames consist of watermark data or not. It helps to discover out all

watermarks contains frame couples. The watermark data can be separated by with certain constraints. In this way, the watermark took away else not. Once the scrambled watermark bit extracted from selected frames than by using a secret key, watermark bits are removed.

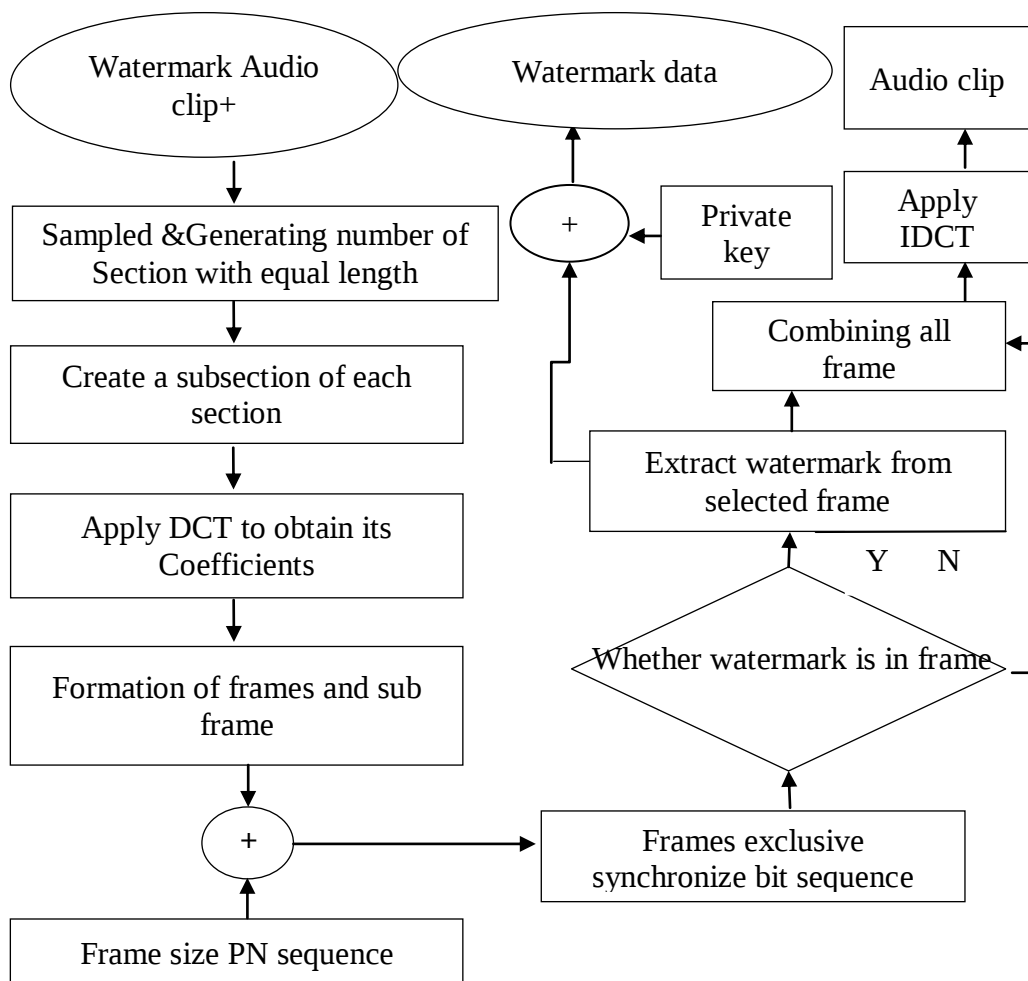


Fig. 2. Watermark removal scheme from watermarked audio clip

III. RESULTS AND ANALYSIS

This section provides simulation illustrations showing the performance of the proposed scheme. We have 10 audio clips that are randomly selected an audio clip in the simulations from different groups such as Bollywood, Pop, Classical, Jazz, and Instrumental.

The length of the audio clip is more than 60 seconds. At 44.1 kHz, samples are section into 16 bits, quantified and collected. A suggested technique for watermarking should be adequate for both conventional and advanced attacks while maintaining high sensory activity efficiency. The suggested watermark's physical property relies on the watermarking parameters ∂, ϕ . Depending on the carefully chosen parametric quantity, the implanting

potential of the suggested scheme can be projected at upto 12 bps. The imperceptibility confirmed by using parameters signal to noise ratio, which is required to a greater extent than 20 dB. Our proposed system offers the average SNR is more than 29dB for chosen audio clips.

This validation evaluates the robustness of the watermark design as compared to conventional and forward attacks. With the following description, the robustness is calculated:

A. Detection rate (DR): The detection rate is calculated between the original watermark and recovered watermark

$$DR = \left(1 - \frac{\text{Number of watermarks bits extracted correctly}}{\text{Number of watermarks bits embedded}} \right) \times 100 \quad (18)$$

TABLE I. AVERAGE PERFORMANCE PARAMETER DR FOR PROPOSED SCHEME FOR DIFFERENT SIGNAL

Attacks	Host audio signal	DR%
Close-loop	Bollywood	97
	Jazz	98
	Instrumental	97
	Pop	98
	Classical	98
Re-sampling	Bollywood	91
	Jazz	90
	Instrumental	80
	Pop	91
	Classical	92
Low filtering pass	Bollywood	98
	Jazz	98
	Instrumental	99
	Pop	97
	Classical	98
High filtering pass	Bollywood	93
	Jazz	97
	Instrumental	93
	Pop	93
	Classical	98
Amplitude	Bollywood	98
	Jazz	97

Pitch scaling	Instrumental	98
	Pop	98
	Classical	97
	Bollywood	93
	Jazz	90
Time scaling	Instrumental	91
	Pop	91
	Classical	92
	Bollywood	91
	Jazz	90
Jitter	Instrumental	81
	Pop	86
	Classical	91
	Bollywood	90
	Jazz	87
Noise	Instrumental	90
	Pop	91
	Classical	92
	Bollywood	80
	Jazz	79
Cropping	Instrumental	90
	Pop	89
	Classical	94
	Bollywood	78
	Jazz	79
MP3	Instrumental	82
	Pop	83
	Classical	80
	Bollywood	98
	Jazz	98
AAC	Instrumental	97
	Pop	98
	Classical	98
	Bollywood	97
	Jazz	98

B. Normalized Cross-Correlation (NCC): The normalized cross-correlation for initial watermark and watermark recovered is determined using the following equation

$$NCC = \frac{\sum_{i=1}^M \sum_{j=1}^N (w(i,j) w'(i,j))}{\sqrt{\sum_{i=1}^M \sum_{j=1}^N (w(i,j))^2} \sqrt{\sum_{i=1}^M \sum_{j=1}^N (w'(i,j))^2}} \times 100 \quad (19)$$

NCC reveals the watermark's quality. Its value lies in the range 0 to 1.

The subsequent traditional and advance attacks are used for the robustness analysis on the watermarked signal:

No attack or Closed loop: The watermarks are removed without any attack.

Re-sampling attack: Down sampling and the up-sampling process are implemented with frequency 16 kHz to 44.1 kHz.

Low-pass filtering (LPF): Filtering with 8 and 12 kHz limit frequency is implemented under this process.

High-pass filtering (HPF): Filtering with 50 - 100 Hz limit frequency is implemented under this process.

Amplitude attack: The amplitudes increased by 1.2 times.

Pitch scaling attack: A Shift in pitches of the watermarked audio signals by using scaling factors 80%.

Time scaling attack: shift in times of the watermarked audio signals by using scaling factors 80%.

Jitter attack: Some samples are randomly detached from the watermarked signals 5000 samples.

Noise attack: Signal to noise (SNR) ratios are of 20 dB is added in the watermark signal.

Cropping attack: From the audio watermarked signals, 100 samples are taken out.

Re-quantization attack: The re-quantization of every sample is done from 16 to 8 bits.

MP3 attack: MPEG 1 Layer III Operation is executed.

AAC attack: Operation of MPEG 4 audio coding system is executed.

The investigational results shown in Table-I suggest that, due to the unconventional signal processing attack, the watermark not considerably pretended by typical attacks but is slightly affected. Even if the low-frequency coefficients are altered, a watermark for a high-energy audio signal remains imperceptible. There is no significant data loss in these regions due to the perceptual importance of elements of relatively low frequency, and watermarked constancy can be enhanced. So, we acquire an optimum value between payload, robustness, and imperceptibility for the precise value of $\gamma = 0.21$ and $\mu = 0.1$. If we exceed this value, then all of the parameters will be slightly influenced. Increases in robustness and weight, and imperceptibility decreases, and improvements in the number of bit errors. The average normalized cross-correlation parameter for all selected signals observed the value is 97.40 %.

IV. CONCLUSION

The proposed method is useful and accurate, for a watermarking scheme of audio signals for audio signal segments using the patchwork approach, which integrates watermark data based on changing the DCT coefficient. For suitable DCT frame pairs, the proposed scheme includes watermarks only to make them highly imperceptible. The mathematical model used for the implanting process with certain chosen DCT frame pairs, whereas the similar insertion is as it is used for searching the watermarked blocks in the decryption step. The design for a suggested strategy with preferred frequency bands and different watermark encryption to ensure a high robustness level. Furthermore, the new strategy is safer due to the need to use a protected key within the decryption process. Results obtained for traditional and advance attacks from simulation demonstrate the better performance of the proposed algorithm.

V. REFERENCES

- [1] G. Hua, J. Huang, Y. Q. Shi, et al., "Twenty Years of Digital Audio Watermarking - A Comprehensive Review." *Signal Processing*, 128: 222–242., 2016.
- [2] D. Kannan and M. Gobi, "An Extensive Research on Robust Digital Image Watermarking Techniques: A Review," *International Journal of Signal Imaging System Engineering*, 8(1/2): 89-104., 2015.
- [3] M. Nosrati, R. Karimi, et al., "Audio Steganography: A Survey on Recent Approaches," *World Appl. Program.*, 2(3): 202–205, 2012.
- [4] A. Valizadeh and Z. J. Wang, "An Improved Multiplicative Spread Spectrum Embedding Scheme For Data Hiding," *IEEE Trans. Inf. Forensics Secur.*, 7(4): 1127-1143, 2012.
- [5] W. Bender, D. Gruhl, et al., "Techniques for Data Hiding," *IBM System Journal*, vol. 35, no. 3/4, pp. 313-336, 1996.
- [6] L. Boney, A. H. Tewfik and K. N. Hamdy, "Digital Watermarks For Audio Signals," In *Proceedings of the Third IEEE International Conference on Multimedia Computing and Systems*, pp. 473- 480, 1966.
- [7] M. Arnold, "Audio Watermarking: Features, Applications, and Algorithms," In *Proceedings of 2000 IEEE International Conference on Multimedia and Expo*, pp. 1013–1016, 2010.
- [8] In-Kwon Yeo and H. J. Kim, "Modified Patchwork Algorithm: A Novel Audio Watermarking Scheme," *IEEE Trans. Speech Audio Process.*, 11(4): 381–386, 2003.
- [9] H. Kang, K. Yamaguchi, et al., "Full-Index-Embedding Patchwork Algorithm for Audio Watermarking," *IEICE Trans. Inf. System*, E91-D (11): 2731–2734, 2008.
- [10] N. K. Kalantari, M. A. Akhaee, et al., "Robust Multiplicative Patchwork Method for Audio Watermarking," *IEEE Trans. Audio, Speech Lang. Processing*, 17(6): 1133–1141, 2009.
- [11] C. Pun and J. Jiang, "Adaptive Patchwork Method for Audio Watermarking Based on Neural Network," *International Journal of Digit Content Technology and its Application*, 5(5): 84–94, 2011.
- [12] I. Natgunanathan, Y. Xiang, et al., "Robust Patchwork - Based Embedding and Decoding Scheme for Digital Audio Watermarking," *IEEE Trans. Audio, Speech Lang. Process.*, 20(8): 2232–2239, 2012.
- [13] I. Natgunanathan, Y. Xiang, et al., "Analysis of A Patchwork - Based Audio Watermarking Scheme," In *Proceedings IEEE ICIEA 2013 Conference*, pp. 900–905, 2013.
- [14] I. Natgunanathan, Y. Xiang, et al., "Robust Patchwork-Based Watermarking Method for Stereo Audio Signals," *Multimedia Tools Appl.*, 72(2): 1387–1410, 2014.
- [15] P. Hu, Q. Yan, L. Dong, et al., "An Improved Patchwork-Based Digital Audio Watermarking in CQT Domain," In *22nd European Signal Processing Conference (EUSIPCO) Proceedings*, pp. 4–7, 2014.