

Secure Loss less Speech Signal Encryption using SNMS, Logistic Chaotic map and Random Sequences

Suresh N. Nakum^a, Mehul B. Shah^{b,*}

^aResearch Scholar Gujarat Technological University, Government Engineering College,
Kankot Road, Rajkot, 360005, Gujarat, India

^bCharutar Vidhyamandal University, G. H. Patel College of Engineering,
Bakrol Road, Anand, 388001, Gujarat, India

ARTICLE INFO

ABSTRACT

Submitted: 1 Dec 2024

Received: 15 Jan 2025

Accepted: 25 May 2025

This paper suggests a lossless way to encrypt and decrypt speech using an algorithm based on a chaotic map. The primary goal is to preserve audio signal confidentiality while preserving bit-level accuracy during decryption. To maintain full audio resolution, the original audio streams are first transformed to a 16-bit integer representation. A Suresh Nakum and Mehul Shah's (SNMS) map and logistic map is used to construct a random chaotic sequence. This sequence is used for two layers of encryption, permutation of audio samples and XOR-based diffusion with 16-bit keys formed from chaotic values. The beginning parameters, which operate as secret keys, are used to construct the exact same chaotic sequence during decryption, and the encryption operations are reversed in the exact order they were performed. The suggested method achieves flawless reconstruction with no quality loss by guaranteeing that the original and decrypted audio signals are bit-exact. The efficacy and resilience of the approach are demonstrated by experimental validation. The encrypted speech signal does not reveal any meaningful information when analyzed through its histogram, frequency spectrum, and spectrogram also the largest difference between the original and decrypted samples is zero. For applications like voice communications, military audio or biometric systems, this method is promising.

Keywords: liyapunov exponent, Suresh Nakum and Mehul Shah's (SNMS) Chaotic Map, Spectrogram, Pseudo Random Number Sequence, Keyspace.

INTRODUCTION

Protecting speech signals during transmission has become essential in the age of fast digital communication, especially in delicate fields like telemedicine, military communications, and intelligent Internet of Things applications. Speech signals present special difficulties for real-time encryption, in contrast to text data. Strong mathematical security is provided by traditional cryptographic approaches like AES, DES and RSA, but because of their high computational complexity and resource requirements particularly in embedded and low-power systems they are frequently inappropriate for real-time audio transmission [1][2].

Researchers have looked into chaotic systems for speech encryption as a solution to these issues. Because of their extreme sensitivity to initial conditions, unpredictability, and ergodicity, chaotic maps are great options for producing safe, pseudo-random sequences that are necessary for encryption [3]. In terms of unpredictability and key sensitivity [4][5], for instance, the application of logistic chaotic maps to scrambling voice signals has demonstrated encouraging outcomes [6].

In addition to chaotic approaches, new neural network based security techniques have been made possible by the development of deep learning, particularly Convolutional Neural Networks (CNNs). Voice-based authentication systems and speaker identification have successfully employed using CNN in [7].

In order to permute voice data over various quantization levels, additional research has suggested multi-map techniques, such as mixing logistic, tent, baker and Bernoulli maps [8][9]. By improving diffusion and confusion qualities, these techniques improve security nevertheless; they may also cause processing overhead and synchronization issues. Some researchers have used Verilog or VHDL to develop chaos-based systems on FPGAs in order to optimize for hardware efficiency [1][10][11], and they have achieved reductions in hardware resource usage of almost $10\times$ [12]. High-dimensional and fractional-order chaotic systems have also been used to increase the complexity of chaotic behavior, which broadens the key space and strengthens defenses against cryptanalysis [13], [14]. These methods, however, have problems with precise synchronization and numerical instability, especially when used on hardware platforms.

Another area of speech related research combines signal transforms such as the Descrete Sine transform (DST), Discrete Cosine Transform (DCT) alone[15]. Fast Fourier Transform (FFT), with chaotic encryption. These systems use chaotic sequences to permute and substitute the real and imaginary parts of the speech spectrum, producing strong encryption results and a decrypted signal with high perceptual quality [2]. To further improve security and bandwidth usage, compressed sensing (CS) has been used to jumble speech signals while lowering transmission overhead [16].

Despite these developments, current methods frequently have shortcomings such as high computational complexity, limited hardware flexibility, tiny key spaces, and susceptibility to differential attacks. These problems underscore the increasing demand for highly secure, hardware-efficient, and lightweight voice encryption techniques that function well in real-time situations. In response, the current work suggests a unique encryption framework that ensures minimal latency and excellent perceptual quality while utilizing the advantages of chaotic systems, particularly those that are tuned for embedded contexts. The goal of the solution is to close the gap between robust security assurances and workable viability for speech communication systems in the real world.

OBJECTIVES

This paper aims to design, implement, and evaluate speech encryption method using advanced chaotic systems for secure voice communication. The proposed approach generates highly sensitive and unpredictable pseudo-random sequences from chaotic maps, forming the core of a robust encryption mechanism. By leveraging chaos theory, the technique ensures strong confusion and diffusion properties essential for cryptographic security. To evaluate the effectiveness and randomness of the encrypted signals, the system incorporates security metrics such as entropy analysis, Number of Samples Changing Rate (NSCR), Unified Averaged Changed Intensity (UACI), and Peak Signal-to-Noise Ratio (PSNR). A primary objective is to eliminate residual intelligibility in the ciphertext, making unauthorized interpretation infeasible. Despite strong encryption, the method preserves speech intelligibility and fidelity upon decryption, maintaining low distortion and high perceptual quality. The technique is computationally efficient, suitable for real-time applications. Large key space enhances resistance to brute-force attacks, while correlation analysis and high mean square error (MSE) values confirm the system's robustness. Additionally, resilience against statistical and differential attacks demonstrates its cryptanalytic strength. Overall, this study presents a secure and efficient scalable voice encryption framework.

METHODS

This encryption methodology leverages the unpredictability and sensitivity of chaotic systems along with the pseudo-randomness of Pseudo Random Noise (PN) sequences to achieve robust and secure signal encryption. The system is designed for signals such as audio, where the integrity and confidentiality of the content are critical.

1. Signal Partitioning

Let the input signal be denoted as a vector $S = [s_1, s_2, \dots, s_N]$

To enable parallel and differential encryption, the signal is divided into four equal sub vectors $S = S_1 \| S_2 \| S_3 \| S_4$

Where each s_i represents a sub-vector of S , with length of $\|S_1\| = \|S_2\| = \|S_3\| = \|S_4\| = N/4$ assuming N is divisible by 4.

2. Chaotic Sequence Generation

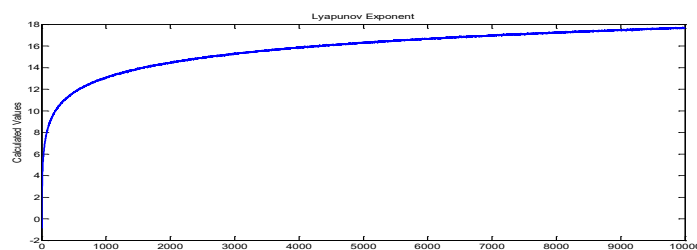


Fig.1. A plot of the Lyapunov exponent illustrating the chaotic behavior of the SNMS map

The chaotic sequences C_1 generated by Suresh Nakum and Mehul Shah's (SNMS) map is a novel designed one dimensional chaotic map having properties of the non-linear dynamic system, which is defined by equation(1). Figure 1 displays the SNMS chaotic map's Lyapunov exponent plot. It is evident from the graph that the Lyapunov exponent is a non-decreasing function, supporting the map's dynamic sensitivity and potential applications in random number generation systems [17].

$$X_{n+1} = r^2 * \sin(\pi X_n^2 - 4) - (1-r) * \cos(\pi X_n^2 - 3) \text{ Eq. --- (1)} \quad \text{and}$$

C_2 are generated using the logistic map equation (2) a well-known chaotic function [2][11]

$$X_{n+1} = r * X_n(1 - X_n) \text{ Eq. --- (2)}$$

Where $X_n \in (0, 1)$, $r \in (3.57, 4]$ ensures chaotic behavior, Initial values $x_0(1)$ and $x_0(2)$ are used to generate C_1 and C_2 , respectively

The sequences are discretized and scaled using equation (3) to match the signal domain range, 16-bit integers, typically by

$$C_k[n] = \lfloor x[n]^{(k)} \times 65535 \rfloor \bmod 65535, k=1, 2. \text{ Eq. (3)}$$

3. PN Sequence and Key Generation

PN sequences Key1, Key2, Key3, and Key4 are generated using linear feedback shift registers (LFSRs) or derived from random seeds. Each sequence matches the length of the respective segment s_i

$$K_i = [key_{i,1}, key_{i,2}, \dots, key_{i,i}, key_{i,N/4}]. \text{ Where } i=1, 2, 3, 4$$

Optionally, four encryption keys Key1 to Key4 can be derived from the chaotic sequences or external parameters, enhancing security.

4. Segment-wise XOR Encryption

Each signal segment is encrypted using an XOR operation with its corresponding PN sequence, as defined in Equation (4).

$$E_i = S_i \oplus K_i \text{ Eq. --- (4)}$$

$i=1, 2, 3, 4$ this operation introduces randomness and diffusion to each sub vector of the signal.

5. Encrypted Signal Concatenation

All four encrypted segments are concatenated

$E = E_1 || E_2 || E_3 || E_4$. Where E is the intermediate encrypted signal of the same length as S .

6. Permutation Using Chaotic Sequence

The vector E is permuted based on a permutation index π derived from chaotic sequence C_1 .

Let: $\pi: \{1, 2, \dots, N\} \rightarrow \{1, 2, \dots, N\}$ be a bijection obtained by sorting C_1 and mapping indices. Equation (5) gives the permuted signal.

$$E' = [E_{\pi(1)}, E_{\pi(2)} \dots E_{\pi(N)}] \text{ Eq. --- (5)}$$

This step introduces confusion and destroys spatial correlation.

7. XOR with Chaotic Sequences

The permuted signal E' undergoes two XOR operations as in equation (6).

$$E'' = (E' \oplus C1) \oplus C2 \text{ Eq. --- (6)}$$

Here, $C1$ and $C2$ are aligned with E' in length. This step increases entropy and nonlinearity.

8. Final Multi-Layer XOR

The output E'' is encrypted through four successive XOR operations equation (7) using the PN sequences.

$$E_{\text{final}} = (((E'' \oplus K1) \oplus K2) \oplus K3) \oplus K4 \text{ Eq. --- (7)}$$

Each layer adds complexity, making reverse engineering or differential analysis significantly more difficult.

The final encrypted output E_{final} is highly nonlinear, sensitive to initial conditions, and exhibits no visible correlation with the original signal S . The scheme is lightweight, and suitable for real-time secure multimedia transmission.

When assessing the performance of a speech signal encryption system, several metrics are used to ensure both security and efficiency [1][17]. One such metric is the PSNR equation (8), which evaluates the level of distortion between the original and decrypted signals. It is computed as

$$\text{PSNR} = 10 \cdot \log_{10} ((A_{\text{max}})^2 / \text{MSE}) \text{ Eq. --- (8)}$$

Where A_{max} is the maximum amplitude of the signal. A lower PSNR indicates more effective encryption. This works in conjunction with the MSE equation (9), which measures the average squared difference between the original and decrypted speech samples

$$\text{MSE} = 1/L \sum [s(k) - s'(k)]^2 \text{ Eq. --- (9) Where } K \text{ range from } 1 \text{ to } L.$$

where $s(k)$ and $s'(k)$ represent the original and encrypted signal samples, and L is the total number of samples. For analyzing resistance to differential attacks, we use two important parameters NSCR equation (10) and UACI equation (11) is defined as

$$\text{NSCR} = 1/L \sum D(k) \times 100 \text{ Eq. --- (10)}$$

Where $D(k) = 1$ if $s'(k) \neq s(k)$ and otherwise. Here, $s(k)$ is the original sample, and $s'(k)$ is the encrypted sample. NSCR quantifies the percentage of changed samples between the original and encrypted signals. Meanwhile, UACI is given by

$$\text{UACI} = 1/L \sum |s(k) - s'(k)| / A_{\text{max}} \times 100 \text{ Eq. --- (11)}$$

That measures the average amplitude variation caused by encryption. A high UACI indicates strong diffusion.

In addition, Shannon Entropy equation (12), calculated as

$$H = -\sum S_p(k) \log S_p(k) \text{ Eq. --- (12)}$$

Where $S_p(k)$ is the probability of sample amplitude level k , assesses the randomness in the encrypted signal, ideally approaching 16 for 16-bit signals.

Finally, Encryption time, measured in milliseconds, determines speed, and key space, ideally larger than 2^{100} , ensures resistance against brute-force attacks. Together, these metrics validate the encryption scheme's robustness and efficiency for secure speech communication.

RESULTS AND DISCUSSION

The audio signals Fig.2 for experimentation were taken from references [18], [19], and [20]. Signal-1 Music [20] was sourced from GetSampleFiles.com and contains WAV audio music files. Signal-2 Military Radio [18] was taken from SampleFocus.com, featuring speech samples related to radio communication. Signal-3 English-Speaking [19] was obtained from the ITU technical report, including English speech data. The encrypted signals have dense, noise-like distributions and no distinguishable speech characteristics, in contrast to the original signal's obvious harmonic or periodic structure Fig.3.

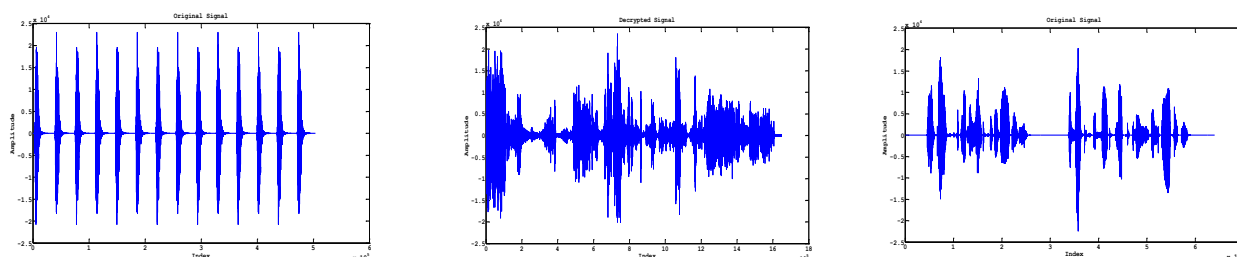


Fig.2. Original Audio Signals 1 2 and 3[18-20]

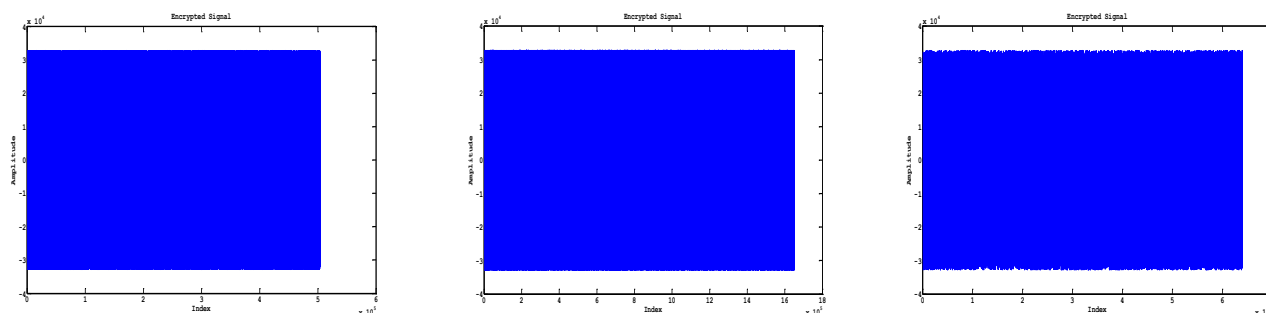


Fig.3. Encrypted Audio Signals1 2 and 3.

The spectral content hides the original information by appearing dispersed over all frequencies Fig.4, which is a desirable feature. Both in the temporal and frequency domains Fig.5, this noise-like appearance guarantees that the encrypted signal closely mimics random noise, rendering it incomprehensible to analytical instruments and human listeners alike. All of these findings support the idea that the encryption method is suitable for high-security audio and voice communication systems since it not only mathematically safeguards the signal but also obfuscates it visually and acoustically.

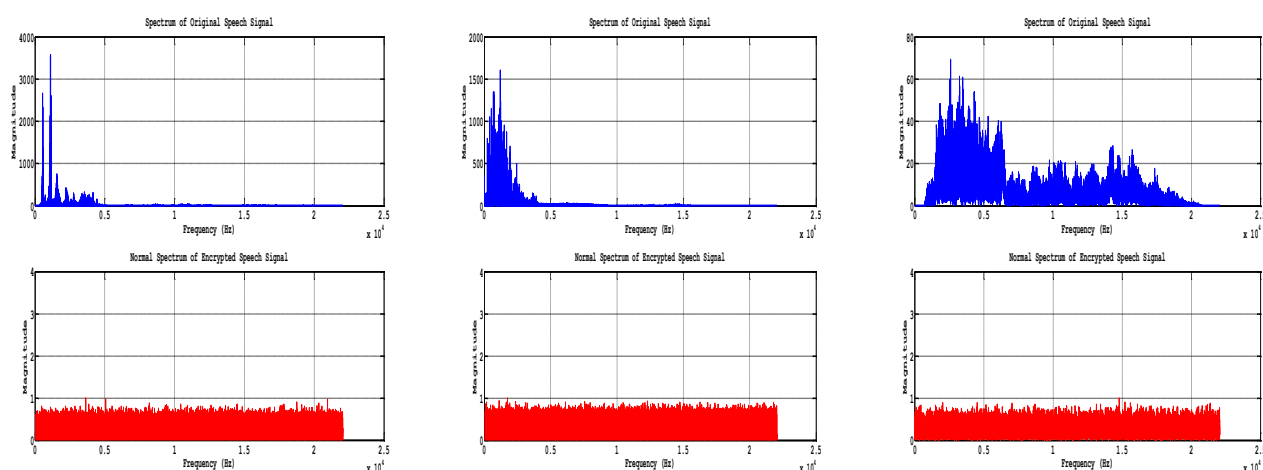


Fig.4. Spectrum of original and Encrypted Audio Signals 1 2 and 3

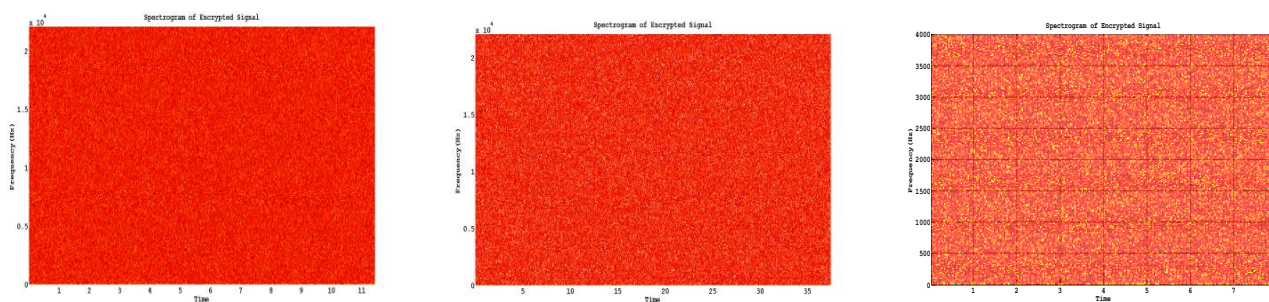


Fig.5. Spectrogram of Encrypted Audio Signals 1 2 and 3

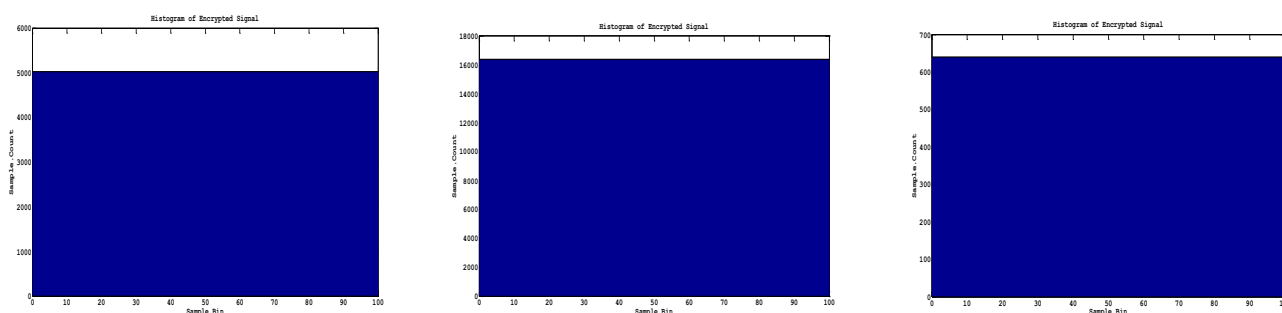


Fig.6. Histogram of Encrypted Audio Signals 1 2 and 3

The enormous key space of 2^{256} , which provides high number of possible key combinations compared to [7] and effectively makes brute-force attacks impossible, strengthens the suggested encryption method's strength. Unauthorized decryption attempts are nearly impossible with this degree of security, even with enormous processing power. Additionally, the robustness of the technique is supported by visual analysis of the encrypted signals. The distribution of amplitude values in the encrypted signals is almost uniform, with no obvious pattern or structure, as can be shown in the histogram plots Fig.6. A characteristic of well-encrypted data is this flat, uniform histogram, which indicates that the output is statistically random.

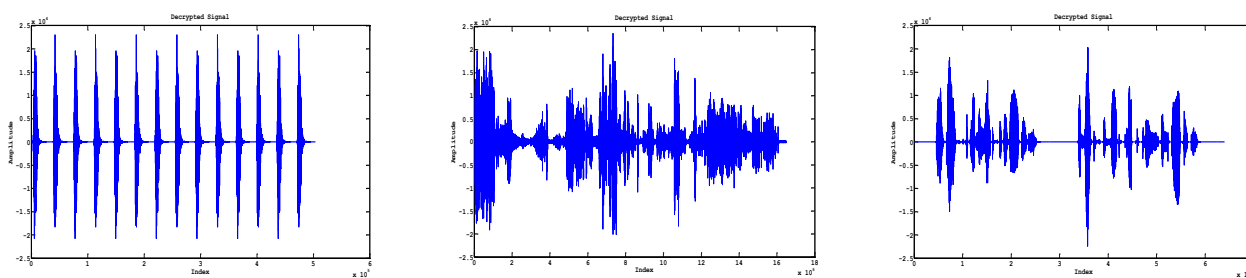


Fig.7. Decrypted Audio Signals 1 2 and 3

Decryption is successful, and the original data was accurately recovered Fig.7. The performance of the suggested encryption technique when applied to three distinct signals is shown and examined in this section. Signal length, encryption time, entropy, correlation, MSE, NSCR, UACI, SNR and PSNR are among the metrics that are used in the evaluation. Table 1 provides a summary of the findings.

Three signals of different lengths were used to assess the encryption method's performance and flexibility. With 503,808 samples in length, the first signal took roughly 0.0559 seconds to encrypt. The third and shortest signal, which was only 64,000 samples long, was encrypted in a quick 0.0079 seconds, while the second, which was the longest at 1,650,176 samples, took about 0.1741. These findings demonstrate that the encryption approach is effective even for huge datasets and scales well with signal size.

Table. 1. Experimentation results for 3 Audio Signals 1 2 and 3

Parameters	Signal-1	Signal-2	Signal-3
	Music[20]	Military Radio[18]	English-Speaking[19]
Length-> L	503808	1650176	64000
MSE	551811407.3432	538569506.8547	540294040.0552
NSCR	99.9999%	99.9999%	99.9999%
Correlation	0.0011	0.0003	-0.0054
UACI	25.35	25.04%	25.10%
PSNR	8.9116 dB	9.0170 dB	9.0032
Entropy	14.9037	14.9717	14.15
SNR	40.3200 dB	36.8595 dB	24.0205
Encryption time->Et	0.055963	0.174140	0.007873
Frame Encryption time= $E_t \cdot S / L = 100u$ Sec Approx., S=no. of samples	98uSec	93uSec	108uSec

The algorithm showed remarkable speed in frame-level processing. Each frame of encryption took about 100 microseconds to complete; Signal-1 took 98 μ s, Signal-2 took 93 μ s, and Signal-3 took 108 μ s. Because of its quick execution, the approach is a good fit for contexts with limited resources or real-time requirements. It proved then [15][13].

The encryption demonstrated strong unpredictability and resilience to statistical attacks by achieving near-maximum entropy. Signal-2: 14.9717. High sensitivity was demonstrated by the NSCR of 99.9999%; in secure systems, almost all samples changed as anticipated. The UACI readings of 25.35%, 25.04%, and 25.10% show steady but noteworthy variations in intensity. A modest correlations 0.0011, 0.0003, -0.0054 that displayed values that were almost zero or negative better compared to [2][4][11][15].

Strong distortion was further confirmed by a high MSE 538–551 million in comparison to [1][10][12]. Secure scrambling is supported by low PSNR ~9 dB compared. 50 dB in [11][2] and decreasing SNR 40.32, 36.86, 24.02 dB. For safe speech encryption, the technique works well and is technically sound.

CONCLUSION

The encrypted signal exhibits a uniform histogram, indicating the absence of amplitude patterns, while the spectrogram shows no identifiable time-frequency features. Similarly, the frequency spectrum lacks distinct peaks, closely resembling random noise. These characteristics confirm that the encryption method effectively conceals all perceptual and statistical features of the original speech signal. The results further demonstrate that the technique offers complete reversibility, along with strong security properties such as high entropy, full sample alteration, and low correlation with the original signal. Notably, it achieves these features with minimal encryption time, making it a practical solution for the secure transmission and storage of voice data.

References:

- [1] M. F. Tolba, H. Saleh, Y. Al Salami, M. Al-Qutayri, and B. Mohammad, "DS2B: Dynamic and secure substitution box for efficient speech encryption engine," *IEEE Access*, vol. 9, pp. 92791–92802, Jul. 2021, doi: 10.1109/ACCESS.2021.3093247.
- [2] P. Sathiyamurthi and S. Ramakrishnan, "Speech encryption algorithm using FFT and 3D-Lorenz–logistic chaotic map," *Multimedia Tools and Applications*, vol. 79, no. 13, pp. 17817–17835, Feb. 2020, doi: 10.1007/s11042-020-08729-5.
- [3] F. F. J. Farsana and K. Gopakumar, "A novel approach for speech encryption: Zaslavsky map as pseudo random number generator," *Procedia Comput. Sci.*, vol. 93, pp. 816–823, 2016.
- [4] Y. Huang, L. Wang, Z. Li, and Q. Zhang, "A new 3D robust chaotic mapping and its application to speech encryption," *Chaos, Solitons & Fractals*, vol. 184, Art. No. 115038, May 2024, doi: 10.1016/j.chaos.2024.115038.
- [5] S. Hashemi, M. A. Pourmina, S. Mobayen, and M. R. Alagheband, "Multiuser wireless speech encryption using synchronized chaotic systems," *International Journal of Speech Technology*, vol. 24, no. 3, pp. 651–663, Mar. 2021, doi: 10.1007/s10772-021-09821-3.

- [6] S. A. Gebereselassie and B. K. Roy, "Speech encryption algorithm based on two newly designed chaotic maps," *Franklin Open*, vol. 5, article 100055, 2023.
- [7] Q. Zhang, Y. Li, Y. Hu, and X. Zhao, "An encrypted speech retrieval method based on deep perceptual hashing and CNN-BiLSTM," *IEEE Access*, vol. 8, pp. 142067–142078, Aug. 2020, doi: 10.1109/ACCESS.2020.3015876.
- [8] P. Sathiyamurthi and S. Ramakrishnan, "Speech encryption using chaotic shift keying for secured speech communication," *EURASIP Journal on Audio, Speech, and Music Processing*, vol. 2017, no. 20, 2017, doi: 10.1186/s13636-017-0118-0.
- [9] N. F. Soliman, Z. Mostfa, F. E. Abd El-Samie, and M. I. Abdalla, "Performance enhancement of speaker identification systems using speech encryption and cancelable features," *International Journal of Speech Technology*, vol. 20, no. 4, pp. 977–1004, Oct. 2017, doi: 10.1007/s10772-017-9435-z.
- [10] M. F. Tolba, W. S. Sayed, M. E. Fouda, H. Saleh, M. Al-Qutayri, B. Mohammad, and A. G. Radwan, "Digital emulation of a versatile memristor with speech encryption application," *IEEE Access*, vol. 7, pp. 174530–174540, Dec. 2019, doi: 10.1109/ACCESS.2019.2957300.
- [11] D. Herbadji, A. Herbadji, I. Haddad, H. Kahia, A. Belmeguenai, and N. Derouiche, "An enhanced logistic chaotic map based tweakable speech encryption algorithm," *INTEGRATION, the VLSI Journal*, vol. 97, Art. No. 102192, Apr. 2024, doi: 10.1016/j.vlsi.2024.102192.
- [12] A. H. Elsafty, M. F. Tolba, L. A. Said, A. H. Madian, and A. G. Radwan, "Enhanced hardware impl. of a mixed-order nonlinear chaotic system and speech encryption application," *Int. J. Electron. Commun. (AEÜ)*, vol. 125, Art. No. 153347, 2020, doi: 10.1016/j.aeue.2020.153347.
- [13] Y. Huang, C. Li, Z. Li, Q. Zhang, and F. Yang, "Fractal matrix speech encryption algorithm based on fractional order robust chaos," *Appl. Acoust.*, vol. 207, Art. No. 108903, 2024, doi: 10.1016/j.apacoust.2023.108903.
- [14] L. J. Sheu, "A speech encryption using fractional chaotic systems," *Nonlinear Dyn.*, vol. 65, pp. 103–108, 2011, doi: 10.1007/s11071-010-9877-1.
- [15] D. Slimani and F. Merazka, "Encryption of speech signal with multiple secret keys," *Procedia Comput. Sci.*, vol. 128, pp. 79–88, 2018, doi: 10.1016/j.procs.2018.03.011.
- [16] L. Zeng, X. Zhang, L. Chen, Z. Fan, and Y. Wang, "Scrambling-based speech encryption via compressed sensing," *EURASIP Journal on Advances in Signal Processing*, vol. 2012, no. 1, Art. No. 257, 2012.
- [17] S. N. Nakum and M. B. Shah, "*SNMS Chaotic Encryption Protect Medical Images against Statistical and Differential Attacks*," unpublished.
- [18] Y. Huang, C. Li, Z. Li, and Q. Zhang, "Efficient speech encryption algorithm based on three-dimensional quadratic exponential robust chaos," *Appl. Acoust.*, vol. 207, Art. No. 109234, 2023, doi: 10.1016/j.apacoust.2023.109234.
- [19] SampleFocus, "Free Speech samples, sounds, and loops," *SampleFocus.com*. [Online]. Available: <https://samplefocus.com/tag/speech>.
- [20] International Telecommunication Union, "Computer processing, data management and energy perspective," FG-AI4EE D.WG2-02, Technical Report, 2024. [Online]. Available: <https://www.itu.int/myworkspace/#/t-signals/vectors?val=17>
- [21] GetSampleFiles.com, "Sample WAV Audio File Download," 2025. [Online]. Available: <https://getsamplefiles.com/sample-audio-files/wav>