

Protecting confidential verbal communication through active noise cancellation^{*}

Oleksandr Laptiev^{1,*}, Vitalii Savchenko^{2,†}, Alla Kobozeva^{3,†}, and Tetiana Laptieva^{4,†}

¹ Taras Shevchenko National University of Kyiv, 60 Volodymyrska str., 01033 Kyiv, Ukraine

² State University of Information and Communication Technologies, 7 Solomianska str., 03110 Kyiv, Ukraine

³ Odessa National Maritime University, 34 Mechnykova str., 65029 Odesa, Ukraine

⁴ National Technical University Kharkiv Polytechnic Institute, 2 Kyrpychova str., 61002 Kharkiv, Ukraine

Abstract

The article highlights a scientifically based method of protecting confidential verbal communication by hiding speech information using specially formed interference. The purpose of the study is to develop an approach that provides effective masking of the speech signal from technical means of interception. To achieve the goal, we analyzed the frequency parameters of speech signals and the sensitivity of modern microphones used to implement unauthorized access to audio information. Since microphones have a wide frequency range and high sensitivity, there is a need to form interference that would cover this range and have sufficient power for effective masking. The interference is formed in such a way as to be deterministic for the legal recipient and random for a third-party analyzer, which ensures minimal correlation between the intercepted signal and the original speech. This significantly complicates its recognition by artificial intelligence or speech recognition systems. We implement the method by digital signal processing, using noise-like signal generation algorithms, which allows to provide a high level of protection in real time. The effectiveness of the proposed method was verified by modeling in the Python environment, using WAV audio files and signal analysis based on Mel-Frequency Cepstral Coefficients (MFCC), which are one of the most common tools for analyzing speech signals. The results of modeling speech information masking proved the adequacy and practical usefulness of the proposed method. Thus, the developed approach provides reliable protection of speech information from technical leakage, while maintaining its quality and accessibility for authorized users, and can be implemented in practical systems of technical information protection.

Keywords

speech information, verbal communication, speech recognition, digital signal processing, interference, information protection, cybersecurity.

1. Introduction

The relevance of the study lies in the fact that at the current stage of development of information technologies and means of information interception, threats of unauthorized access to confidential verbal communications are increasing. This is due to the widespread use of highly sensitive microphones, as well as modern speech analysis systems, which, based on recorded signals, can automatically recognize and extract the content of a conversation. Traditional methods of protecting speech information, such as acoustic masking or shielding of premises, often do not provide a sufficient level of security, especially when using high-tech interception means. Not all eavesdropping devices can be detected by existing search methods. Remote-controlled radio microphones, wired microphones, passive resonators, mini dictaphones — almost all of these devices are not detected by standard methods. Even a modern mobile phone can contain a built-in dictaphone, which means that any phone lying on the table can be used to record a conversation.

^{*} CSDP'2026: Cyber Security and Data Protection, May 7, 2026, Lviv, Ukraine

^{*} Corresponding author.

[†] These authors contributed equally.

✉ alaptiev64@ukr.net (O. Laptiev); savitan@ukr.net (V. Savchenko); alla_kobozeva@ukr.net (A. Kobozeva); tetiana1986@ukr.net (T. Laptieva)

ORCID 0000-0002-4194-402X (O. Laptiev); 0000-0002-3014-131X (V. Savchenko); 0000-0001-7888-0499 (A. Kobozeva); 0000-0002-5223-9078 (T. Laptieva)



© 2026 Copyright for this paper by its authors. Use permitted under Creative Commons License Attribution 4.0 International (CC BY 4.0).

2. Problem statement

Microphones used in modern devices have a wide frequency range and high sensitivity, which allows you to receive a high-quality signal even at a distance or in conditions of background noise. In addition, speech recognition systems based on signal processing technologies, such as MFCC, significantly increase the efficiency of analyzing intercepted data. In this regard, there is a need to develop new approaches to protecting speech information that would ensure its concealment from technical means of interception while at the same time allowing it to be restored by the legal recipient without loss of quality. Ensuring real-time security is especially important when delays or loss of information are unacceptable, for example, during confidential verbal communications.

Therefore, the purpose of this article is to resolve this contradiction. The contradiction is resolved by the proposed method of adding a specially formed synchronized interference to the useful speech signal. Such interference covers the speech frequency range (300 Hz – 8 kHz), has sufficient power for effective masking, but is reproducible by a legal receiver, ensuring accurate restoration of the original signal. This approach makes it significantly more difficult or impossible to recognize speech information by interception means, while maintaining its availability for authorized users without loss of quality.

3. Related works overview

The issue of protecting speech information by superimposing (adding) the same signal, which is in antiphase to the main one, is considered in many publications. Thus, work [1] is a review article (tutorial) on Active Noise Cancellation (ANC), which provides algorithms (feedforward, feedback, adaptive), implementation methods and limitations that are the foundation for understanding the technical limits of phase cancellation of a speech signal. Publication [2] presents ANC schemes with one sensor, relevant for the case when it is necessary to minimize the hardware complexity when attempting local speech cancellation. The authors of article [3] describe the use of ANC systems for protecting speech messages. They highlight some implementation features, in particular, the possibility of signal processing via electrical or electromagnetic channels for accelerated formation of antiphase even before the sound wave arrives. It is noted that the efficiency is valid up to 4 kHz, which covers the speech spectrum.

In article [4], a new deep neural network approach for Active Speech Cancellation (ASC) is proposed. It uses frequency division to more accurately generate an anti-signal with high phase coherence, using noise and speech for effective suppression. In [5], a system is proposed to protect speech confidentiality by playing background music with a hidden high-frequency signal. This signal can later be removed using an adaptive filter, which reduces the background noise level by more than 20 dB. The authors of [6] describe the use of adaptive filters to suppress a predetermined “voice”. This method describes how an anti-phase component is generated to reduce unwanted sound.

In [7], a system is described with the ability to select an adaptive algorithm based on noise characteristics. The filtering goal is fast convergence and effective noise suppression up to 30 dB. The authors of [8] investigate the application of deep learning for ANC, in which the trained model can selectively suppress noise while passing the speech signal; they also consider the creation of a “quiet zone” in space. The study [9] shows the possibility of suppressing the speech signal using a secondary source located near the speaker’s mouth, which generates a linearly predicted inverse signal. The paper [10] considers methods to make speech unreadable while preserving the structure of the acoustic scene (an approach to privatization of recordings, related to the idea of “destruction” of speech content without total sound deformation). The authors of [11] consider an approach that exploits the nonlinearity of microphones and ultrasonic “jammers” to protect against eavesdropping (an alternative implementation of speech protection).

The annotated report [12] describes the practical implementation of Active Voice Control (AVC) technology for speech masking and discusses the difficulties of its real-world implementation

(delays, feedbacks, spatial effects). The publication [13] presents a study that analyzes the effectiveness of sound masking for privacy protection and demonstrates the weaknesses of existing masking systems. The paper [14] considers low-latency ANC in the context of deep learning (i.e. deep ANC), for which a time-domain method using a recurrent network (ARN) is used, which reduces the algorithmic delay of deep ANC. The paper is useful for real-world systems where latency is critical for generating anti-speech signals. The paper [15] describes alternative, engineering methods for ensuring speech privacy (masking, sound absorption, architectural solutions).

The publications presented show that modern approaches to concealing speech information — from classical active noise cancellation with phase inversion to ultrasonic jammers and deep neural network algorithms — are capable of significantly reducing speech intelligibility for unauthorized microphones and observers, especially in controlled conditions and a limited area. Their advantages lie in the relative technical simplicity of implementation (in digital channels), the ability to adapt to the noise environment and effectiveness against standard interception methods. However, the disadvantages include critical sensitivity to phase and amplitude errors, dependence on room acoustics, delays in the processing path, and the risk of complete loss of information for everyone if the restoring signal or key is absent.

This leads to a key contradiction: on the one hand, it is necessary to ensure a high level of concealment of speech information from technical means of interception, and on the other hand, to preserve the possibility of its accurate restoration by authorized users.

The purpose of this paper is to develop a method for hiding speech information that effectively makes it impossible to recognize it by technical interception means, while ensuring the possibility of accurate signal recovery by the legal recipient. The scientific task is to create interference that is added to the useful speech signal in such a way that its spectrum overlaps the range of human speech (from 300 Hz to 8 kHz), and its power is sufficient for effective masking. In this case, the interference must be synchronized, i.e. reproduced by the legal recipient to restore the original signal, and at the same time look like random noise to an external analyzer. This ensures minimal interconnection between the intercepted signal and the original speech, which significantly complicates its analysis by artificial intelligence or speech recognition systems. Thus, solving this scientific task allows achieving a high level of protection of speech information from technical leakage, while maintaining its quality and ease of use for authorized users.

4. Mathematical foundations of Active Noise Cancellation

Our task is to ensure the concealment (masking) of speech information $s(t)$ by adding a barrier to it $n(t)$, so that the resulting signal $x(t) = s(t) + n(t)$ was difficult for third parties to recognize, but allowed for the possibility of recovery $s(t)$ by legal recipient. To do this, we will formulate the problem in the following form.

Let the useful signal (speech) be as [16]:

$$s(t): R \rightarrow R, t \in [0, T], \quad (1)$$

where $s(t)$ is analog or discrete signal containing speech information.

Speech signal frequency range: $f \in [f_{min}, f_{max}]$, for example, [300, 3400] Hz for telephone speech or [300, 8000] Hz for human speech.

At the same time, the protection of speech information should be considered within a broader information protection framework, where artificial intelligence and blockchain technologies are increasingly used to strengthen confidentiality, integrity, and resilience of cybersecurity systems [17]. Two frequency bands are considered here because the frequency range of voice transmission in telephones (including mobile communications) is usually limited to save bandwidth and ensure quality with efficient coding. This is the so-called telephone band (PSTN — Public Switched Telephone Network), and it is also used in modern mobile networks (2G, 3G, 4G), although new standards may support wider bands (e.g. HD Voice / VoLTE).

The second frequency range is the range of speech that is perceived by the human ear. The human ear can theoretically hear sounds in the range from 20 Hz to 20 kHz, but for speech perception, the main information is in a different range: 300 Hz–8 kHz. Low frequencies (300–1000 Hz) carry more information about vowel sounds; high frequencies (1–8 kHz) are important for consonant sounds, which affects speech intelligibility.

A person can understand speech even in a narrower range (e.g., 300–3400 Hz), but higher quality and naturalness is achieved when transmitting up to 7–8 kHz [16, 18]. This is also consistent with the development of intelligent cyber defence systems, where artificial intelligence technologies are used to improve detection, adaptation, and protection mechanisms in complex information environments [19].

To intercept a conversation, various modifications of secret information acquisition tools are used, but in the vast majority of such devices, the primary one is a microphone. The quality and range of interception of speech information depends on its sensitivity. Therefore, we consider it necessary to investigate the parameters of the primary source of information interception — namely microphones. For the transmission of speech information, it is important that the microphone “catches” the range of human speech well — approximately 300 Hz–8 kHz.

Typical frequency response of speech microphones:

1. Telephony microphones: 300 Hz–3400 Hz.
2. Professional speech recording microphones: 80 Hz–15 kHz.
3. VoIP/headset/smartphone microphones: 100 Hz–10 kHz.

Such microphones have a uniform sensitivity across the speech range to ensure maximum intelligibility. Table 1 shows the sensitivity ranges of microphones by design features [20, 21]. However, confidential verbal communication can also be exposed through optical and laser-based technical leakage channels, where the protective capabilities of glass under laser sounding depend on its elemental composition [22]. In addition, single-layer reflective coating films of hafnium dioxide can be considered as a material-based countermeasure against special cyber espionage devices [23].

Table 1

Examples of microphone sensitivity

Type	Sensitivity	Frequency range
Dynamic	1-5 mV/Pa	50 Hz-15 kHz
Condenser	5-50 mV/Pa	20 Hz-20 kHz
Electret	5-20 mV/Pa	20 Hz-16 kHz
Smartphone microphone	5-15 mV/Pa	100 Hz-10 kHz (optimized for the speech)

Analysis of the data in Table 1 shows that condenser or electret microphones are most often used for high-quality recording of speech information due to their high sensitivity and ability to accurately transmit speech signals.

The next step will be to determine the parameters of the interference, which can be described by the expression:

$$n(t): R \rightarrow R, t \in [0, T], \quad (2)$$

where $n(t)$ is a function that depends on the variable t . In our case, $n(t)$ can represent the noise that is added to the useful signal $s(t)$.

But it should be taken into account that the interference (noise) is generated in such a way that:

- it covers the same frequency range as $s(t)$;
- it has sufficient power to mask $s(t)$;
- it is a deterministic or random function that can be reproduced by a legal receiver.

Then the resulting signal will be written as:

$$x(t) = s(t) + n(t). \quad (3)$$

In our context, $n(t)$ is the noise (interference) added to the desired signal $s(t)$, and both functions $s(t)$ and $n(t)$ are real-valued functions of time: $s(t): R \rightarrow R$, $n(t): R \rightarrow R$. Resulting signal $x(t) = s(t) + n(t)$ will also be a function of time with real values $x(t): R \rightarrow R$. This is a signal that is transmitted over a communication channel and is accessible to a potential attacker [24].

Let's form the requirements for the noise $n(t)$:

- Frequency domain: $Supp(N(f)) \subseteq [f_{min}, f_{max}]$, where $N(f)$ – Fourier transform $n(t)$.
- Interference power: $P_n = \frac{1}{T} \int_0^T n^2(t) dt \geq \alpha P_s$, where P_s is power of the desired signal, $\alpha > 0$ is masking parameter (i. e. $\alpha = 1$).
- Reproducibility (synchronized interference), it must be possible for the legal recipient to recover $n(t)$ exactly or approximately in order to perform compensation: $\hat{s}(t) = x(t) - \hat{n}(t)$, where $\hat{n}(t)$ is assessment of interference at reception.

Recent studies on laser-sounding conditions additionally confirm that the elemental composition of glass influences its protective properties, which is relevant for complex protection of rooms and communication environments where confidential speech is transmitted [25].

An optimization problem arises, which consists in finding a function $n(t)$ that satisfies the constraints [26]:

- $Supp(N(f)) \subseteq [f_{min}, f_{max}]$,
- $P_n \geq \alpha P_s$,
- $\exists$ synchronization/playback algorithm $n(t)$,
- $D(s, \hat{s}) \leq \varepsilon$, where D is measure of distance between signals, ε – precision.

As examples of optimality criteria, it is possible to use, such as:

1. Minimizing signal visibility/audibility: $\max_{n(t)} I(x, s)$, where $I(x, s)$ is mutual information between the transmitted and received signal (should be minimal for the attacker).
2. Maximizing the degree of masking: $\max_{n(t)} \left(\frac{\|s(t)\|}{\|x(t) - n(t)\|} \right)$.

Based on the above, the mathematical formulation of the problem consists in constructing an interference signal $n(t)$ that:

- operates in the same frequency range as speech;
- has sufficient power for masking;
- is recoverable by the legitimate recipient;
- provides a given quality of information concealment.

This task can be implemented using digital signal processing methods, coding theory, noise-like signals, synchronized chaos systems, etc. The method of software implementation proposed by us consists in using digital signal processing methods.

To confirm the adequacy of the proposed method we developed a Python program. As an audio signal to be protected, we used an audio file in the “wav” format. For correct modeling, we will use the Mel-scale – Mel-Frequency Cepstral Coefficients (MFCCs), which is an empirical scale based on the human perception of sound frequency and is currently the leader among known speech recognition systems. The cepstral is one of the types of homomorphic signal processing, a function of the inverse Fourier transform of the logarithm of the signal power spectrum. The cepstral can be written as follows:

$$C_s(q) = \frac{1}{2\pi} \int_0^{2\pi} |S(\omega)|^2 e^{i\omega q} d\omega, \quad (4)$$

where $S(\omega)$ is input signal spectrum; q is an argument that has the dimension of time, but it is cepstral time based on the fact that $C_s(q)$ at any moment q depends on the function $s(t)$ input signal with spectrum $S(\omega)$, given at $-\infty < t < \infty$ [27].

5. Simulation and discussion of results

Fragments of the program code for implementing the proposed method with comments are shown in Figure 1.

```
# 1. Reading audio files
audio_path = 'sample-12s.wav'
noise_path = 'sample-6s.wav'

# Loading the main signal
signal, sr = librosa.load(audio_path, sr=16000)
# Loading interference with the same sampling rate
noise, _ = librosa.load(noise_path, sr=sr)
# Signal length equalization
min_len = min(len(signal), len(noise))
signal = signal[:min_len]
noise = noise[:min_len]
# Signal mixing
mixed_signal = signal + 0.3 * noise
# The coefficient 0.3 adjusts the interference level
# 2. Calculating MFCC for all signals
def compute_mfcc(y):
    return librosa.feature.mfcc(
        y=y, sr=sr,
        n_mfcc=13,
        n_fft=2048,
        hop_length=512 )
mfcc_orig = compute_mfcc(signal)
mfcc_noise = compute_mfcc(noise)
mfcc_mixed = compute_mfcc(mixed_signal)
```

Figure 1: Fragment of the program code for implementing the method for protecting the confidential verbal communications.

The simulation results, namely the graphical representation of the signal, are shown in Figure 2.

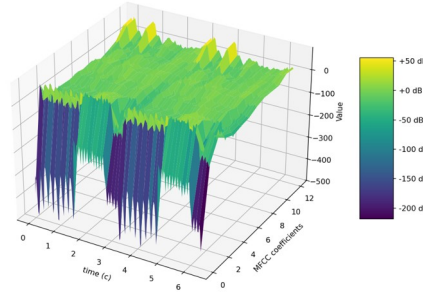


Figure 2: The signal to be protected in Mel-Frequency Cepstral Coefficients (MFCCs) parameters.

The next step for a more visual presentation of the results of modeling the results of protecting confidential verbal communications is to present the interference signal that meets the above conditions, Figure 3.

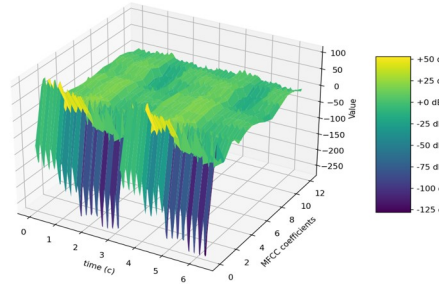


Figure 3: Schematic of the barrier that will be used to protect confidential verbal communications. The final graph showing the results of the protection method is shown in Figure 4.

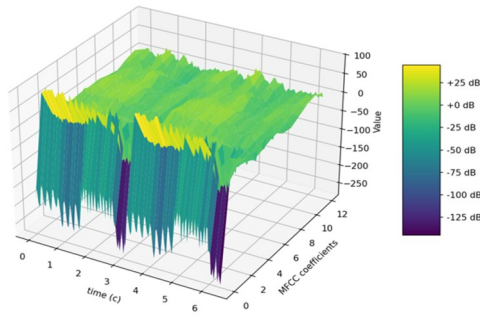


Figure 4: Graph of the resulting signal that can be intercepted or recorded in a room where confidential verbal communications are taking place.

6. Conclusions

In this paper we have proposed a method for ensuring the security of confidential verbal communications based on hiding speech information by adding a specially formed interference to it. This approach makes it difficult to obtain information by unauthorized means of interception, in particular, using microphones, which are the main devices for implementing technical data leakage. The goal is to make the speech signal inaccessible to an outside observer, but to preserve the possibility of its accurate restoration by the legal recipient.

The work analyzes the frequency characteristics of speech signals, in particular, it was found that the main information is contained in the range from 300 Hz to 8 kHz, where low frequencies provide the transmission of vowel sounds, and high frequencies provide the intelligibility of consonants. It is also taken into account that telephone systems use a narrower range — from 300 to 3400 Hz, which is sufficient for basic speech intelligibility, but does not provide maximum quality. To effectively mask speech, the interference must overlap this frequency range and have sufficient power to hide the desired signal.

Special attention is paid to the sensitivity of microphones, since it is on it that the possibility of intercepting speech information depends. It has been established that condenser and electret microphones, which have high sensitivity and a wide frequency range, are the most effective for recording speech. The frequency characteristics of microphones used in modern devices, such as smartphones and VoIP devices, which are optimized for speech, are also considered.

The mathematical model of the information hiding process involves the formation of interference in such a way that it operates in the same frequency range as the useful signal, is synchronized for the legal recipient and provides minimal interconnection with the original signal for a third-party analyzer. This reduces the likelihood of recognizing speech information by technical interception means.

To verify the effectiveness of the method, a simulation was conducted in the Python environment, where an audio file in WAV format was used, and the signal analysis was performed using Mel-Frequency Capsular Coefficients (MFCC), which is an effective tool for analyzing speech signals. Based on the graphs of the signal, interference, and the resulting signal, a comparative analysis was performed, which showed that the resulting signal is significantly different from the original, which confirms the effectiveness of the proposed method.

Thus, the developed method for ensuring the security of confidential verbal communications allows you to effectively hide speech information from technical means of interception, while maintaining the possibility of its recovery by the legal recipient. This is achieved by forming and adding synchronized interference that meets certain frequency and power requirements. The method is implemented using modern methods of digital signal processing, which makes it promising for practical application in speech information protection systems.

General summary: The work offers a scientifically sound method for protecting speech information from unauthorized receipt, which can be implemented in real conditions to ensure the confidentiality of verbal communications.

Declaration on Generative AI

The authors have not employed any Generative AI tools.

References

- [1] S. Kuo, D. Morgan, Active noise control: A tutorial review, *Proc. IEEE*, 87(6) (1999), 943-973. doi:10.1109/5.763310.
- [2] A. Oppenheim, E. Weinstein, K. Zangi, M. Feder, D. Gauger, Single-sensor active noise cancellation, in *IEEE Trans. Speech Audio Process.*, 2(2) (1994), 285-290. doi:10.1109/89.279277.
- [3] S. Lizunov, E. Filobok, Protecting Speech Information Using Active Sound Canceling Systems 23(1) (2021), 20-25. doi:10.18372/2410-7840.23.14959.
- [4] Y. Mishaly, L. Wolf, E. Nachmani, Deep Active Speech Cancellation with Mamba-Masking Network, 2025. doi:10.48550/arXiv.2502.01185.
- [5] J. Huang, A. Ito, Study on the Background Music Cancellation System for Speech Privacy. In 2021 6th International Conference on Signal and Image Processing, ICSIP-2021, 811-815. doi:10.1109/ICSIP52628.2021.9688835.
- [6] D. Venkata Ram Reddy, A. Prabhu, M. Sri Krishna Chaitanya, K. Gopiram, Active Noise Cancellation for Predefined Voices, *Int. J. Eng. Technol.* 7(2.7) (2018), 897-899. doi:10.14419/ijet.v7i2.7.11090.
- [7] R. Ramli, A. Noor, S. Abdul Samad, Noise cancellation using selectable adaptive algorithm for speech in variable noise environment, *Int. J. Speech Technol.* 20 (2017), 535-542. doi:10.1007/s10772-017-9425-1.
- [8] H. Zhang, D. Wang, Deep ANC: A deep learning approach to active noise control, *Neural Networks*, 141 (2021) 1-10. doi:10.1016/j.neunet.2021.03.037.
- [9] K. Kondo, K. Nakagawa, Speech emission control using active cancellation, *Speech Commun.*, 49(9) (2007), 687-696. doi:10.1016/j.specom.2007.04.010.
- [10] M. TAILLEUR, M. Lagrange, P. Aumond, V. Tourre, Enforcing Speech Content Privacy in Environmental Sound Recordings using Segment-wise Waveform Reversal. (2025). doi:10.48550/arXiv.2507.08412.
- [11] M. Gao, et al., Cancelling Speech Signals for Speech Privacy Protection against Microphone Eavesdropping. In *Proceedings of the 29th Annual International Conference on Mobile Computing and Networking (ACM MobiCom '23)*. Assoc. Comput. Mach., New York, USA, 2023, 36, 1-16. doi:10.1145/3570361.3592502.
- [12] C. Rose. Active Voice Control: An Implementation of Active Noise Control for Canceling Speech. Thesis Abstract for Bachelor of Electrical Engineering, Auburn University Honors College, Alabama, USA, 2007, p. 54.
- [13] S. Anand, P. Walker, N. Saxena, Compromising Speech Privacy under Continuous Masking in Personal Spaces, 2019 17th International Conference on Privacy, Security and Trust (PST), Fredericton, NB, Canada, 2019, 1-10, doi:10.1109/PST47121.2019.8949053.
- [14] H. Zhang, A. Pandey, D. Wang, Low-Latency Active Noise Control Using Attentive Recurrent Network. *IEEE/ACM Trans. Audio, Speech, Lang. Process.* 31, 2023, 1114-1123. doi:10.1109/taslp.2023.3244528.
- [15] J. Bruyninckx, Tuning the office sound masking and the architectonics of office work, *Sound Stud.*, 9(1) (2023), 64-84. doi:10.1080/20551940.2022.2162765.
- [16] O. Laptiev, et al., The method of spectral analysis of the determination of random digital signals, *Int. J. Commun. Networks Inf. Secur. (IJCNIS)* 13(2) (2021), 271-277. doi:10.54039/ijcnis.v13i2.5008.
- [17] O. Harasymchuk, I. Opirskyy, et al. Mod. Methods Ensuring Inf. Prot. Cybersecur. Syst. Using Artif. Intell. Blockchain Technol. Kharkiv: Technology Center PC, p. 132, (2025). doi:10.15587/978-617-8360-12-2.

- [18] O. Laptiev, et al., Algorithm for Recognition of Network Traffic Anomalies Based on Artificial Intelligence. In: 5th International Congress on Human-Computer Interaction, Optim. Robot. Appl. (HORA), 2023. doi:10.1109/HORA58378.2023.10156702.
- [19] O. Harasymchuk, I. Opirskyy, et al., *Intell. Cyber Def. Syst. Detect. Ransomware Prot. Wirel. Networks Based Artif. Intell. Technol.*, Kharkiv: Technology Center PC, p. 182, (2025). doi:10.15587/978-617-8360-22-1.
- [20] ETSI ES 202 057-1,-2,-3,-4: Speech Processing, Transmission and Quality Aspects (STQ); User related QoS parameter definitions and measurements, 2021.
- [21] ISO/IEC 27001:2022, Information security, cybersecurity and privacy protection. Inf. secur. manag. syst. Requir.
- [22] N. Dzianyi, et al., Investigation of the Protective Capabilities of Glass from Laser Sounding Depending on Its Elemental Composition. *EUREKA: Phys. Eng.*, 2022 (5), pp. 162-174. doi:10.21303/2461-4262.2022.002527.
- [23] V. Dudykevych, et al., Investigation of the Use of Protective Characteristics of Single-Layer Reflective Coating Films of Hafnium Dioxide to Counter Special Cyber Espionage Devices. In: *Proceedings of the 11th IEEE International Conference on Intelligent Data Acquisition and Advanced Computing Systems: Technology and Applications, IDAACS 2021*, vol. 2, pp. 709-713. doi:10.1109/IDAACS53288.2021.9660890.
- [24] N. Lukova-Chuiko, et al., The method detection of radio signals by estimating the parameters signals of eversible Gaussian propagation. In: *IEEE 3rd International Conference on Advanced Trends in Information Theory*, 2021, 67-70. doi: 10.1109/ATIT54053.2021.9678856.
- [25] B. Korchak, N. Dzianyi, I. Opirskyy, M. Shved, Determining the Influence of Glass Elemental Composition on Its Protective Properties Using Laser-Sounding. *Chem. Chem. Technol.*, 19(4), pp. 751-760, (2025). doi:10.23939/chcht19.04.751.
- [26] M. Devaiah, Protection of Confidentiality and Security of Parties in Online Dispute Resolution, *Int. J. Multidiscip. Res.*, 7(2) 2025. doi:10.36948/ijfmr.2025.v07i02.39369.
- [27] S. Davis, P. Mermelstein, Comparison of Parametric Representations for Monosyllabic Word Recognition in Continuously Spoken Sentences. *IEEE Trans. Acoust., Speech, and Signal Processing*, 28(4) (1980) 357-366. doi:10.1109/TASSP.1980.1163420.