

Pilot Study to Enhance Cover-Selection-Based Audio Steganography (CAS) Using Feed-Forward Neural Network

Taqiyuddin Anas¹, Farida Ridzuan^{2*}, Sakinah Ali Pitchay³

^{1,3}Faculty of Science and Technology, Universiti Sains Islam Malaysia, 71800, Nilai, Malaysia

²Cybersecurity and Systems Research Unit, Faculty of Science and Technology, Universiti Sains Islam Malaysia, 71800 Nilai, Malaysia

ARTICLE INFO

Article history:

Received 24 October 2024

Revised 29 December 2024

Accepted 10 January 2025

Published 1 March 2025

Keywords:

Cover Selection

Carrier Selection

Steganography

Audio Steganography

Feed-forward Neural Network

Machine Learning

DOI:

10.24191/jcrinn.v10i1.490

ABSTRACT

Steganography is a method of concealing a hidden message inside another medium ranging from image to video. The specification of the cover audio used for message embedding plays a role in the whole steganography performance. The Cover-Selection-Based Audio Steganography (CAS) technique addressed cover selection in audio steganography. However, finding the optimal cover audio using the CAS technique currently takes a significant amount of time. Therefore, the CAS technique is improved by utilising a machine learning technique called Feed-Forward Neural Network (FFNN). Similarly to CAS, Least Significant Bit (LSB) encoding is utilised for data embedding. The proposed technique's effectiveness is assessed by comparing it with CAS regarding time performance, precision, and the stego audio quality, using a dataset of 95 inputs. The pilot study demonstrated that the FFNN model achieved 60% precision over the CAS technique in machine learning evaluation. For the audio stego evaluation, the finding shows that the proposed technique performed slightly lower than the CAS technique in the imperceptibility aspect while performing better than the CAS technique in the robustness and capacity aspects. The proposed technique achieved faster cover selection with a 5,126.89% speed reduction in performance evaluation. This study offers a valuable reference for future research on audio steganography, particularly in enhancing the performance of cover selection using machine learning.

1. INTRODUCTION

Steganography is the art and science of hiding information in plain sight. It is a technique that allows for the concealment of a message, image, or file within another message, image, or file, making it difficult to detect the presence of the hidden information. The word "steganography" came from the Greek words "steganos", meaning "covered or concealed", and "graphein", meaning "to write" (Fridrich, 2011). The existing literature on cover selection steganography has primarily focused on the relationship between medium characteristics and steganography performance (Ren-e, 2014), and the impact of cover selection

^{2*} Corresponding author. E-mail address: farida@usim.edu.my
<https://doi.org/10.24191/jcrinn.v10i1.490>

on steganalysis (Wang et al., 2019). To determine the effectiveness of any audio-steganographic approach, three common audio-steganography characteristics must be evaluated which is capacity, imperceptibility and robustness (Bhowal et al., 2017). Capacity refers to the amount of hidden information embedded within the cover message. Imperceptibility refers to how well a hidden message is embedded in the cover audio without affecting the audio and robustness refers to the ability of a hidden message to withstand attacks or compression.

Most research in audio steganography has been conducted on message-embedding steganography methods to produce a good quality stego-embedded. However, there is a gap in the research in terms of time complexity analysis. Moreover, the recent work on cover selection implementations in audio steganography, which is the Cover-Selection-Based Audio Steganography (CAS) technique, has some limitations. The primary concern with CAS is the significant amount of time it takes to select the optimal cover audio (Noor Azam, 2023). Noor Azam's study focused on the balance of robustness, imperceptibility, capacity, and security of audio stego characteristics without emphasizing its time performance efficiency. It is also computationally costly and time-consuming, which is not viable for real-time applications.

To allow the implementation of cover selection in the real-time application of audio steganography, this study will improve the performance of time in selecting the optimal cover for audio steganography based on the CAS technique of Noor Azam (2023) supervised Feed-Forward Neural Network (FFNN).

2. LITERATURE

There are a few carriers or cover files of hidden messages in steganography, which are text, image, audio, and video (Artz, 2001). Each cover file has a slightly different method of encoding messages. Steganography carrier choice depends on communication context, hidden message size or type, and security level. This research focused on audio for higher security and resilience against detection and attacks.

Cover selection refers to identifying and selecting suitable audio files, known as cover files, that can be used to embed hidden messages. The selection of an appropriate cover is crucial in steganography, as it directly impacts the undetectability and security of the hidden information (Ren-e, 2014; Wang et al., 2020). Existing cover selection techniques mainly were on images (Amin Seyyedi & Ivanov, 2014; Andono & Setiadi, 2023; Bin Li et al., 2015; Subhedar, 2021; Wang et al., 2019, 2020; S. Wu et al., 2015; Yuan & Chen, 2014).

Few audio cover selection techniques have been proposed, with the latest by Noor Azam et al. (2023) called the Cover-Selection-Based Audio Steganography (CAS) technique. These methods vary in strengths, focusing on imperceptibility, capacity, or robust security. However, few address time efficiency. The CAS technique optimizes the trade-off between imperceptibility, capacity, and robustness using a multi-objective evolutionary algorithm (MOEA) for cover audio selection but overlooks time performance. For data embedding, the Block-based Chaotic Multi-level LSB (BCM-LSB) method was used as it offers superior dynamic security compared to Rashid (2020) (Noor Azam et al., 2023). The general framework of CAS is shown in Fig. 1. This research aims to close that gap by enhancing time efficiency using the Feed-Forward Neural Network (FFNN) algorithm.

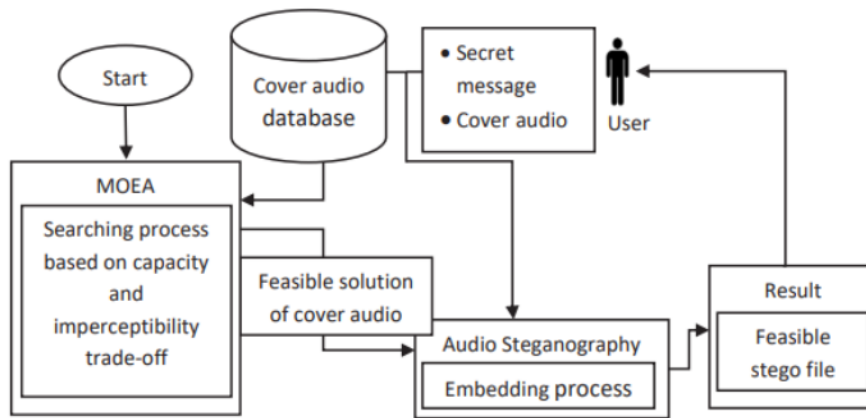


Fig. 1. General framework of cover audio selection

Source: Noor Azam et al. (2023)

3. METHODOLOGY

Selecting the most appropriate cover audio file requires the identification of patterns, characteristics, and relationships between the cover audio files and the hidden message, which can be challenging to achieve. Hence, FFNN was chosen for this research due to its ability to learn patterns and relationships in data. The FFNN model produced in this study will be utilized to determine which audio cover is suitable for a required hidden message, hence optimizing time efficiency without compromising stego quality. The FFNN implementation flow can be seen in Fig. 2.

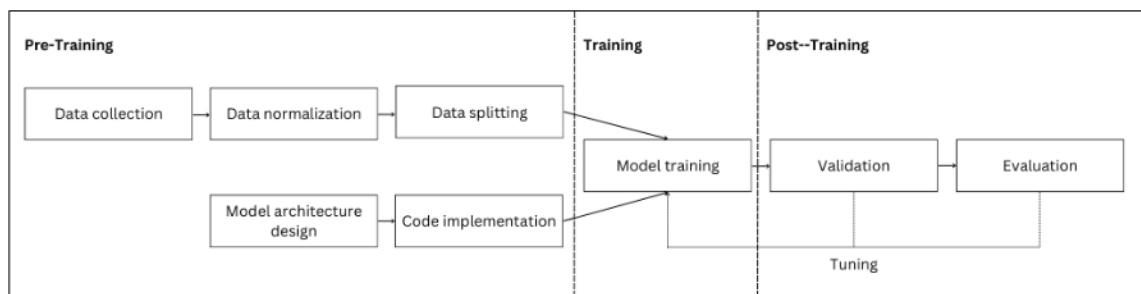


Fig. 2. Implementation flow of feed-forward neural network

Based on Fig. 2, the implementation flow is divided into three phases, which are pre-training, training, and post-training.

3.1 Pre-training

The pre-training phase consists of two sections, which are data collection and architecture design, which are closely related and crucial for establishing a basis for the FFNN model in this study.

3.1.1 Data Collection

Data collection, data normalization, and data splitting are the three main steps in the pre-training phase to produce the required datasets. Audios that were utilized for pre-training consisted of stego audio and the original audio, along with the output of evaluation data that were taken from the CAS procedure. In the realm of cover-selection techniques, the approach by Sajedi & Jamzad (2009) is closely aligned with this research objective. The methodology incorporates crucial elements that resonate with this study, notably by utilising cover audio and stego audio outputs as input for FFNN models. Data normalization is then carried out, including cleaning, transformation, and standardizing data to maintain dataset quality.

3.1.2 Model Architecture Design

The architecture design of the FFNN defines how the model learns from the data. FFNN network structure was split into two parts, training network and post-training network. Network structures and code implementations for model creation were inspired by Wu et al. (2016). The training network structure is shown in Figure 3, while the selection network structure is laid out in Figure 4.

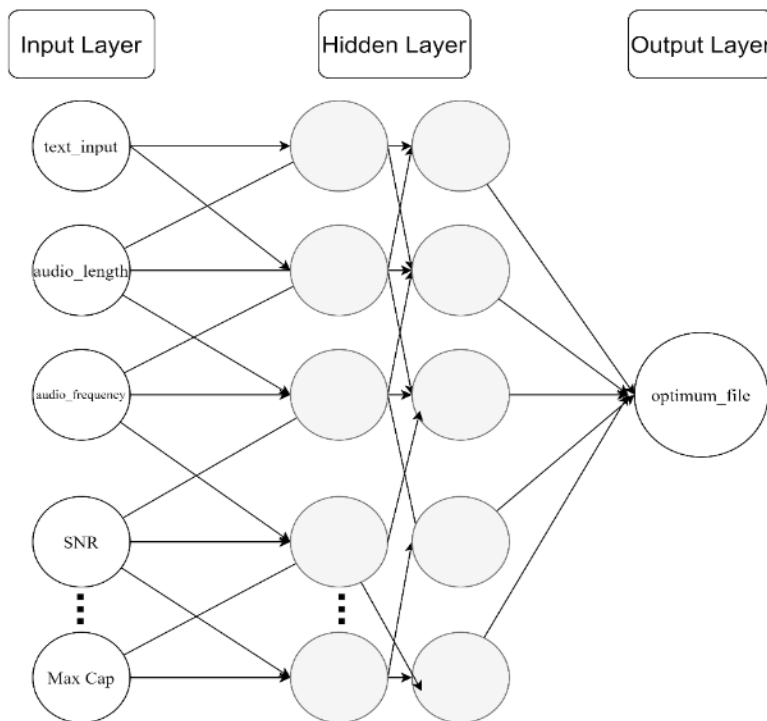


Fig. 3. Training Network structure with multiple features as input

FFNN training network structure is shown in Fig. 3, where datasets generated during the pre-training phase were used as input during the training phase. These features, generated during the pre-training phase, serve as inputs for the hidden layer, where the network learns patterns in the hidden layer along with the optimum_file from the same dataset.

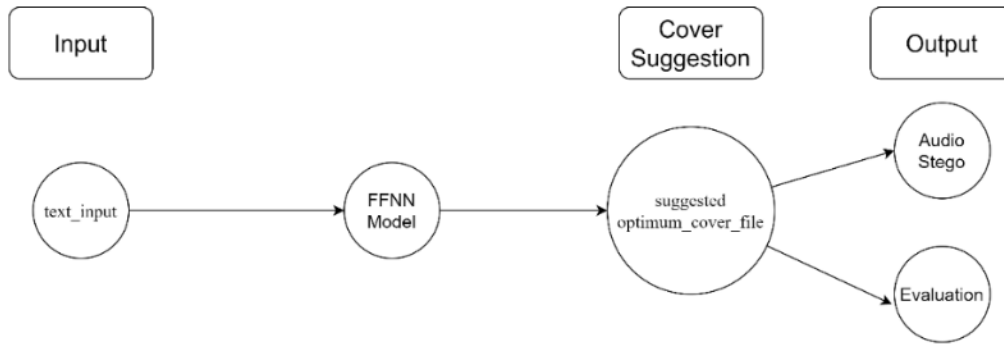


Fig. 4. Cover Selection network structure

Fig. 4 illustrates the cover selection network structure, showcasing the process of identifying the most suitable cover audio. This structure ensures optimal matching for input data and cover audio by leveraging a trained FFNN model for precise selection.

3.2 Training

In the training phase, the training network was trained using datasets produced from Phase 1. The model training utilised machine learning parameters as specified in Table 1.

Table 1. Parameter configuration for FFNN training

Parameter	Value
numInputs	1
numLayers	2
numOutputs	1
numInputDelays	0
numLayerDelays	0
numFeedbackDelays	0
numWeightElements	91
sampleTime	1

MATLAB configuration outlines a neural network setup with specific parameters, as in Table 1. The network was designed with one input layer, two hidden layers, and one output with weight elements 91. The initial FFNN training network structure took multiple input features following findings by Ridzuan et al. (2024) for audio steganography cover selection.

The number of inputs was set to one due to the input being a combined array of various features. Features input that used were text_input, optimum_file, audio_length, audio_frequency, PSNR, SNR, MSE, BER and Max Cap. In particular, the input was defined as 'input = [text_input; optimum_file; audio_length; audio_frequency; PSNR; SNR; MSE; BER,Max Cap]'. The training pseudocode is shown in Fig. 5. Instead of having multiple separate inputs, all relevant features were concatenated and fed into the network per row. The result of this phase, which is the FFNN model, was then used in the post-training phase.

```

% Initialize and clear environment
clc
clear

% Load data from CSV file
datacell = readtable('normalized_output_uat.csv');

% Extract relevant columns from the dataset
text_input = datacell.text_input';
optimum_file = datacell.optimum_file';
audio_length = datacell.audio_length';
audio_frequency = datacell.audio_frequency';
PSNR = datacell.PSNR';
SNR = datacell.SNR';
MSE = datacell.MSE';
BER = datacell.BER';

% Prepare input array for training
input = [text_input; optimum_file; audio_length; audio_frequency; PSNR;
        SNR; MSE; BER];

% Initialize feedforward neural network with 2 hidden layers
net = feedforwardnet(2);

% Train the network using input and target output (optimum_file)
net = train(net, input, optimum_file);
save net.mat

% Display network structure and trained network
disp(net)
view(net)

```

Fig. 5. Pseudocode of Training Network implementation with multiple features as input to produce FFNN model, net.mat

While iteration of training was performed, the FFNN model was validated using k-fold cross-validation since the dataset was less than 100 and limited (Vabalas et al., 2019). The datasets were divided into k subsets, and the model was trained on k-1 folds while validated on the remaining fold in each iteration. This process was repeated for all folds, ensuring every data point was used for training and validation.

3.3 Post-training

In post-training phase, the trained model was imported and utilised to choose an optimum cover that was suitable for the hidden message using the selection network established in Fig. 4. The implementation pseudocode of the selection network is presented in Fig. 6.

```

% Initialize: Clear all variables, console, and figures
clc;
clear all; % Clear variables and workspace

% Load the trained FFNN model from the saved file
load('net.mat'); % Load trained network

% Define the input data to simulate the network
input_afterTrain = 'Lorem Ipsum'; % Example input data

% Simulate the network using the provided input to get the output
output_afterTrain = sim(net, input_afterTrain); % Simulate FFNN output

% Display the simulated output
disp(output_afterTrain); % Show output from the network

```

Fig. 6. Pseudocode of selection network implementation with text as input and suitable cover as output

Based on Figure 6, the cover selection process occurs prior to any data injection into the audio file, in contrast to the CAS technique, which iterates the entire injection and audio steganography evaluation

process with each submission of a hidden message. This iterative approach in the CAS technique causes significant timing execution issues.

4. RESULTS AND DISCUSSIONS

This section presents the analysis of the performance of the proposed technique. Similar evaluation input was used for both the CAS technique and the proposed technique, and the results were recorded. Three evaluations were conducted, which are (i) machine learning model evaluation, (ii) audio stego characteristic evaluation, and (iii) time performance evaluation.

4.1 Machine Learning Evaluation

For the machine learning model evaluation, the model was trained with 95 data samples. The results were compared against the CAS technique. In this research, only precision was used to measure the performance of the machine-learning model. The result of the model evaluation is shown in Table 2.

Table 2. Result of FFNN model evaluation

Metrics	Description	Count
True Positive (TP)	The model correctly identified the audio cover as suitable or better for embedding a hidden message compared to the CAS technique	57
False Positive (FP)	The model incorrectly classified a non-optimal audio cover as suitable for steganography	38
Total		95

This study is primarily interested in the precision of predictions when choosing steganographic audio covers. The result of the precision percentage is shown in Table 3.

Table 3. Calculation of FFNN model precision.

Metrics	Description	Formula	Result
Precision	Percentage of the closeness of the outcome to each other	$\frac{TP}{TP + FP} \times 100\%$	60%

Based on Table 3, the result of the FFNN model indicates a precision of 60% against the CAS technique, showing the model's moderate success in selecting suitable audio covers. The precision is above average compared to the implementation by Deep Learning, which achieved 49% precision (Das et al., 2023). However, this suggests room for further optimisation in enhancing the performance of the FFNN model for steganographic purposes by having more extensive datasets of 1,000 inputs similar to Ye et al. (2019) implementation.

4.2 Audio Stego Characteristic Evaluation

For steganography evaluation, audio stego is evaluated through robustness, imperceptibility, and capacity.

4.2.1 Imperceptibility

Imperceptibility evaluation begins with metrics such as Signal-to-Noise Ratio (SNR), Peak Signal-to-Noise Ratio (PSNR), and Mean Square Error (MSE). Each of these measures has different aspects of performance, but all focus on evaluating the imperceptibility of the steganographic method. SNR measures

the ratio between the signal strength and the background noise, assessing the clarity of the stego audio after data embedding. The formula used for SNR is:

$$\text{Signal Noise Ratio} = 10 \times \log_{10}((\text{signal_power})/(\text{noise_power})) \quad (1)$$

where “signal_power” is the power of the original cover, and “noise_power” is the power of the noise present in the audio stego (embedded) signal. The result of SNR can be seen in Fig. 7.

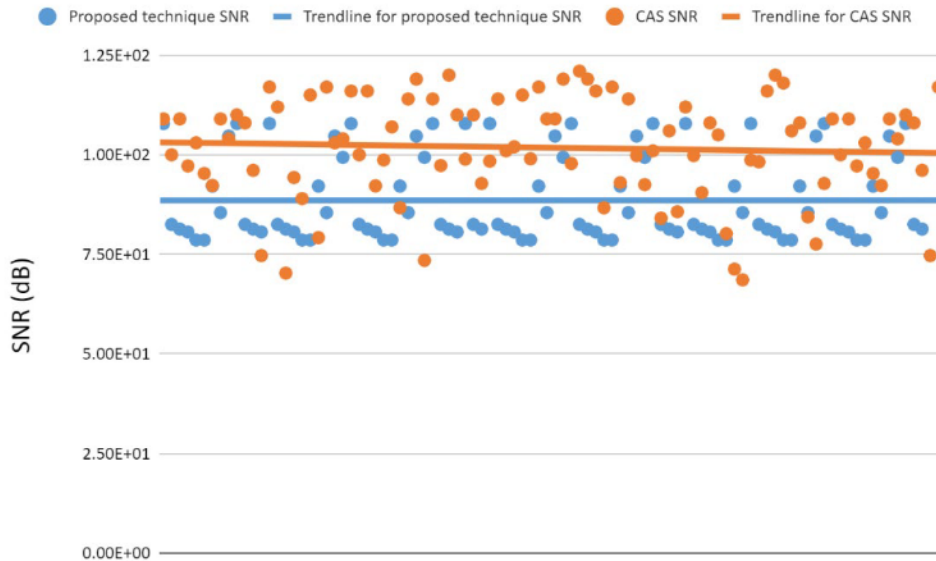


Fig. 7. Comparison of SNR value between proposed technique and CAS

Based on Figure 7, it is observed that the proposed technique exhibits greater consistency and precision compared to the CAS technique. The standard deviation for the proposed technique is 10.67, while for the CAS technique, it is 12.86, indicating that the CAS technique has more scattered data and is less precise. However, the average SNR value for the proposed technique is 88.59 dB, slightly lower than the CAS technique's average of 101.34 dB. This may be due to a lack of data compared to the well-established Noor Azam (2023) technique that utilizes different value usage of bps as a trade-off between capacity and imperceptibility, whereas the proposed method uses it as constant 1 bps.

In addition, PSNR is commonly used to assess the quality of steganographic methods by measuring the peak error between the original and the stego audio. PSNR is particularly important in determining how much the hidden data distorts the cover audio, with higher values indicating better imperceptibility and lower distortion. The formula used for PSNR is:

$$\text{Peak Signal Noise Ratio} = 10 \times \log_{10}\left(\frac{MAX^2}{MSE}\right) \quad (2)$$

where MAX refers to the maximum possible pixel value of the image, and MSE is the Mean Squared Error between the original and stego signals. PSNR result comparison between CAS and the proposed technique can be seen in Fig. 8.

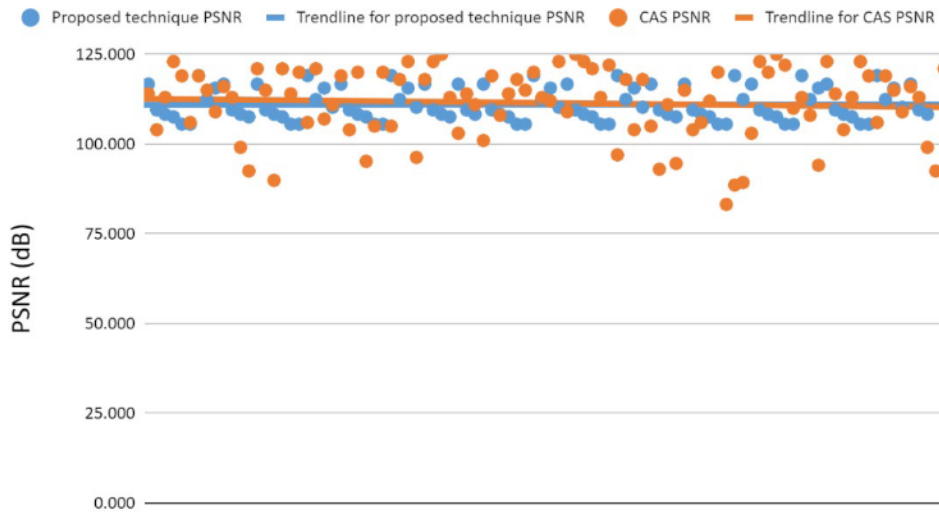


Fig. 8. Comparison of PSNR value between the proposed technique and CAS technique

Based on Figure 8, the PSNR analysis shows that the proposed technique and CAS techniques achieve similar imperceptibility, with average PSNR values of 110.992 dB for the proposed technique and 111.314 dB for CAS. However, the proposed technique demonstrates significantly better consistency, as indicated by its lower standard deviation (4.435) compared to CAS (9.899). This suggests that while both methods provide comparable audio quality, the proposed is more reliable and precise, with less variation in performance, whereas CAS exhibits more inconsistent results.

4.2.2 Robustness

Robustness is evaluated with Bit Error Rate (BER), which measures the ratio of incorrectly received bits to the total number of bits transmitted, providing insight into the stego's resilience against errors during transmission or manipulation. A lower BER indicates a more robust steganographic technique capable of preserving the hidden data, even in noise or attacks. BER formula is:

$$\text{Bit Error Rate} = \frac{\text{Number of bit errors}}{\text{Total number of bits}} \quad (3)$$

The results of the BER evaluation is presented in Fig. 9.

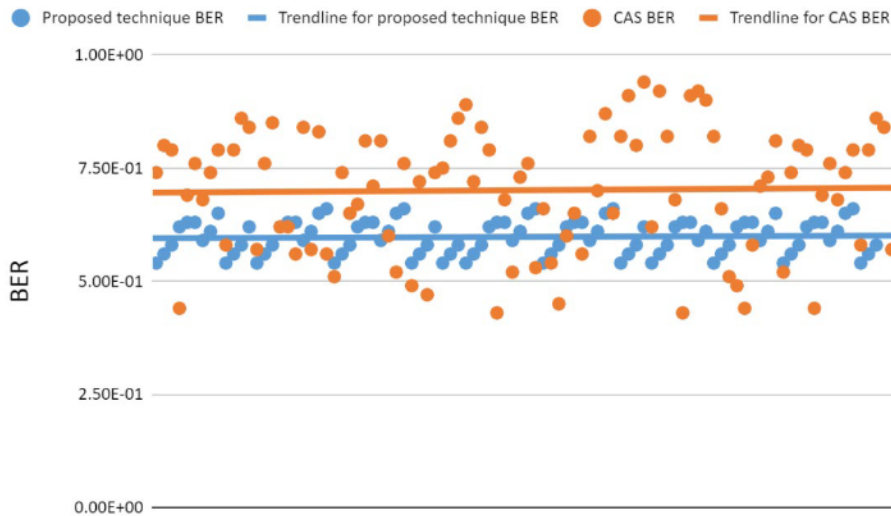


Fig. 9. Comparison of BER value between the proposed technique and CAS technique

Based on Fig. 9, the BER results show that the proposed technique significantly outperforms the CAS technique in precision. The average BER for the proposed technique is $5.98\text{E-}01$, lower than CAS, which has an average BER of $6.98\text{E-}01$. The proposed technique offers more accurate message extraction with fewer bit errors. Additionally, the standard deviation of BER for the proposed technique is 0.0368 , much smaller than the 0.1367 standard deviation for CAS. This highlights that the proposed technique produces more consistent results, while CAS exhibits more significant variability in its performance

4.2.3 Capacity

The capacity evaluation utilizes the Maximum capacity instead of the Embedding rate. The formula for maximum capacity per audio sample is as follows:

$$\text{Max Capacity (bits)} = \text{Sampling Rate} \times \text{Duration(seconds)} \times \text{Number of Channels} \times \text{bps} \quad (4)$$

where the Sampling Rate is a constant of 44100Hz , Duration varies per cover used, the Number of Channels is a constant of 1 due to mono audio and bit per second (bps) is set as 1 bps. The result is presented in Fig. 10.

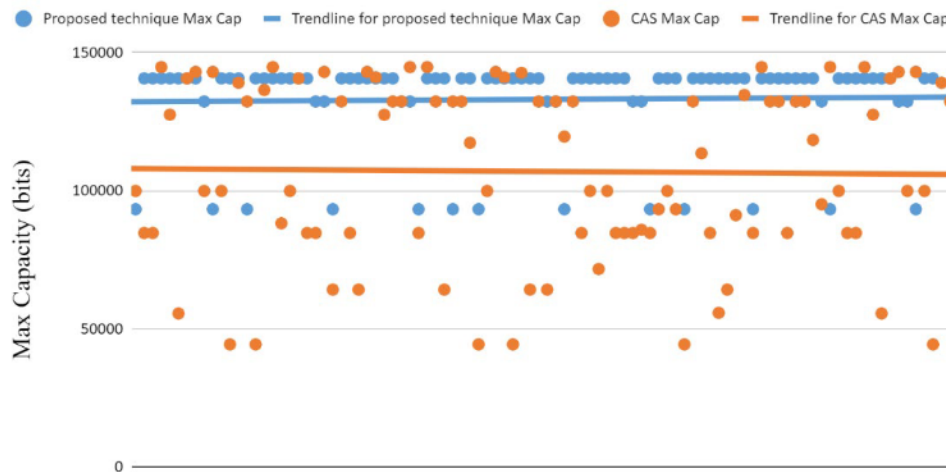


Fig. 10. Comparison of Max Capacity value between the proposed technique and CAS technique

Based on Fig. 10, maximum capacity results show that the proposed technique outperforms the CAS technique with an average capacity of 133,477.61 bits compared to 107,315.80 bits for CAS. Additionally, the proposed technique demonstrates more consistency with a lower standard deviation of 15,865.78, while CAS exhibits more significant variability with a higher standard deviation of 31,958.75. This indicates that the proposed technique not only offers higher data embedding capacity but also provides more reliable and stable performance, making it a more effective and consistent approach for audio steganography.

4.3 Time Performance Evaluation

The comparison of timing execution between CAS and the proposed technique shows a significant improvement in time efficiency. CAS required 136470.755 seconds for processing ten inputs, averaging a high execution time per input. In contrast, the proposed technique requires 26.61 seconds for ten inputs. The result of execution per input is calculated as in Table 4.

Based on Table 4, the percentage of reduced timing is 5,126.89%, highlighting the efficiency gains provided by the proposed method. This significant difference can be attributed to the machine learning approach used in the proposed method, specifically the Feed-Forward Neural Network (FFNN), which optimizes the cover selection process. This massive reduction in execution time not only speeds up the steganography process but also makes the proposed technique more scalable and practical for real-time applications. The results suggest that the proposed technique is far more efficient and effective in terms of time performance, making it a clear improvement over the CAS technique.

Table 4. Calculation of time reduction using the proposed technique.

Timing execution per input		Percentage of reduced timing
CAS technique 13,647.07 s	Proposed technique 2.66 s	5,126.89%

5. CONCLUSION AND RECOMMENDATIONS

The evaluation of the Feed-Forward Neural Network (FFNN) model demonstrated its effectiveness in selecting suitable audio covers for steganography. With a precision rate of 60%, the model showed above average in identifying optimal covers. The proposed technique outperformed the CAS technique across a few steganographic characteristics. For imperceptibility, the proposed technique exhibited greater consistency and precision, with a better standard deviation in SNR and PSNR values despite a slightly lower average SNR. Regarding robustness, the proposed technique achieved a lower Bit Error Rate (BER) and demonstrated more consistent results. For capacity, the proposed technique presents a better maximum capacity compared to the CAS technique. Lastly, time performance evaluation revealed a significant improvement in efficiency with the proposed technique, which drastically reduced the execution time against the CAS technique, achieving a 5,126.89% reduction in processing time. This remarkable efficiency gain can be attributed to the machine learning approach used in the proposed technique, specifically the FFNN. The results indicate that the proposed technique is more efficient and practical for real-time applications. In the future, the FFNN model will be improved by using larger datasets. Despite some minor shortcomings in steganography performance, its consistency adds significant value. The proposed technique also offers faster computational time, with only a slight compromise in output quality.

6. ACKNOWLEDGMENTS

This research was funded by the Ministry of Higher Education (MOHE) Malaysia under the Fundamental Research Grant Scheme (FRGS/1/2020/ICT02/USIM/02/1). The authors would like to express their gratitude to Universiti Sains Islam Malaysia (USIM) and MOHE for the support and facilities provided.

7. CONFLICT OF INTEREST STATEMENT

The authors agree that this research was conducted in the absence of any self-benefits, commercial or financial conflicts and declare the absence of conflicting interests with the funders.

8. AUTHORS' CONTRIBUTIONS

Taqiuddin Anas: Conceptualisation, methodology, formal analysis, visualization, investigation, data curation, writing-original draft, references; **Farida Ridzuan:** Conceptualisation, funding acquisition, methodology, supervision, validation and writing – review & editing; **Sakinah Ali Pitchay:** Conceptualisation, supervision, validation and writing – review & editing.

9. REFERENCES

- Amin Seyyedi, S., & Ivanov, N. (2014). Statistical image classification for image steganographic Techniques. *International Journal of Image, Graphics and Signal Processing*, 6(8), 19–24. <https://doi.org/10.5815/ijigsp.2014.08.03>
- Andono, P. N., & Setiadi, D. R. I. M. (2023). Quantization selection based on characteristic of cover image for PVD Steganography to optimize imperceptibility and capacity. *Multimedia Tools and Applications*, 82(3), 3561–3580. <https://doi.org/10.1007/s11042-022-13393-y>
- Artz, D. (2001). Digital steganography: Hiding data within data. *IEEE Internet Computing*, 5(3), 75–80. <https://doi.org/10.1109/4236.935180>
- Bhowal, K., Sarkar, D., Biswas, S., & Sarkar, P. P. (2017). A steganographic approach to hide hidden data <https://doi.org/10.24191/jcrinn.v10i1.490>

- in digital audio based on XOR operands triplet property with high embedding rate and good quality audio. *Turkish Journal Of Electrical Engineering & Computer Sciences*, 25, 2136–2148. <https://doi.org/10.3906/elk-162-267>
- Bin Li, Ming Wang, Xiaolong Li, Shunquan Tan, & Jiwu Huang. (2015). A strategy of clustering modification directions in spatial image steganography. *IEEE Transactions on Information Forensics and Security*, 10(9), 1905–1917. <https://doi.org/10.1109/TIFS.2015.2434600>
- Das, D., Durafe, A., & Patidar, V. (2023). An Efficient Lightweight LSB Steganography with deep learning steganalysis. In M. D Patil, G. K Birajdar, & S. S Chaudhari, *Computational Intelligence in Image and Video Processing* (1st ed., pp. 131–154). Chapman and Hall/CRC. <https://doi.org/10.1201/9781003218111-7>
- Noor Azam, M. H. (2023). Enhancement Of Cover Selection – Based Audio Steganography (CAS) using Block-Based Chaotic Multi – Level Lsb (BCM– LSB) For Balanced Performance [Publication, University Sains Islam Malaysia]. <https://oarep.usim.edu.my/entities/publication/458e87db-1676-4ea8-afcb-020d0c4e5847/full>
- Noor Azam, M. H., Mohd Ridzuan, F. H., & Mohd Sayuti, M. N. S. (2023). Optimized cover selection for audio steganography using multi-objective evolutionary algorithm. *Journal of Information and Communication Technology*, 22(2), 255–282. <https://doi.org/10.32890/jict2023.22.2.5>
- Rashid, R. Dh. (2020). Cover image selection for embedding based on different criteria. In S. S. Agaian, S. P. DelMarco, & V. K. Asari (Eds.), *Mobile Multimedia/Image Processing, Security, and Applications 2020* (p. 30). SPIE. <https://doi.org/10.1117/12.2560720>
- Ren-e, Y. (2014). Cover selection for image steganography based on image characteristics. *Journal of Optoelectronics-laser*. https://www.researchgate.net/publication/288164492_Cover_selection_for_image_steganography_based_on_image_characteristics
- Ridzuan, F., Taqiyuddin Anas, & Sakinah Ali Pitchay. (2024). Cover selection in steganography: A systematic literature review. *Journal of Advanced Research in Applied Sciences and Engineering Technology*, 107–129. <https://doi.org/10.37934/araset.52.2.107129>
- Sajedi, H., & Jamzad, M. (2009). Secure cover selection steganography. In J. H. Park, H.-H. Chen, M. Atiquzzaman, C. Lee, T. Kim, & S.-S. Yeo (Eds.), *Advances in Information Security and Assurance* (Vol. 5576, pp. 317–326). Springer Berlin Heidelberg. https://doi.org/10.1007/978-3-642-02617-1_33
- Subhedar, M. S. (2021). Cover selection technique for secure transform domain image steganography. *Iran Journal of Computer Science*. <https://doi.org/10.1007/S42044-020-00077-9>
- Vabalas, A., Gowen, E., Poliakoff, E., & Casson, A. J. (2019). Machine learning algorithm validation with a limited sample size. *PLOS ONE*, 14(11), e0224365. <https://doi.org/10.1371/journal.pone.0224365>
- Wang, Z., Li, S., & Zhang, X. (2019). Towards improved steganalysis: When Cover selection is used in steganography. *IEEE Access*, 7, 168914–168921. <https://doi.org/10.1109/ACCESS.2019.2955113>
- Wang, Z., Zhang, X., & Qian, Z. (2020). Practical cover selection for steganography. *IEEE Signal Processing Letters*, 27, 71–75. <https://doi.org/10.1109/LSP.2019.2956416>
- Wu, H.-Z., Wang, H.-X., & Shi, Y.-Q. (2016). Can machine learn steganography? - Implementing LSB substitution and matrix coding steganography with feed-forward neural networks (arXiv:1606.05294). arXiv. <http://arxiv.org/abs/1606.05294>

- Wu, S., Liu, Y., Zhong, S., & Liu, Y. (2015). What makes the stego image undetectable? *Proceedings of the 7th International Conference on Internet Multimedia Computing and Service*, 1–6. <https://doi.org/10.1145/2808492.2808539>
- Ye, D., Jiang, S., & Huang, J. (2019). Heard more than heard: An audio steganography method based on GAN. *ArXiv*. <https://www.semanticscholar.org/paper/1c33d4d27728d1b9f4589e04ab5a26cfe961a30>
- Yuan, J., & Chen, H. (2014). *Embedding Suitability Adaptive Cover Selection for Image Steganography*: 2014 International Conference on e-Education, e-Business and Information Management (ICEEIM 2014), Shanghai, China. <https://doi.org/10.2991/iceeim-14.2014.11>



© 2025 by the authors. Submitted for possible open access publication under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

About the Authors

Taqiuddin Anas, BSc is a graduate of Information Security and Assurance Program, Faculty of Science and Technology, Universiti Sains Islam Malaysia. His main research activity is in steganography and machine learning. His career interest related to blue team in cybersecurity, focusing on malware reversing, rules detection, and threat hunting. He can be reached through his email at ahmadtaqiuddin98@raudah.usim.edu.my

Farida Ridzuan, PhD is an Associate Professor in the Information Security and Assurance Program, Faculty of Science and Technology, Universiti Sains Islam Malaysia. She earned her first-class B.Sc. (Hons.) in Computer Science from Universiti Teknologi Malaysia, an M.Sc. in Discrete Mathematics from the University of Essex, U.K., and a Ph.D. from Curtin University, Australia. Her research focuses on steganography and cryptography, with numerous publications in top-tier journals and RM2 million in secured research funding. She can be reached at farida@usim.edu.my.

Sakinah Ali Pitchay, PhD is an Associate Professor in the Information Security and Assurance Program at Universiti Sains Islam Malaysia. She received her PhD in Computer Science from the University of Birmingham, UK, a Master's degree in Software Engineering from Universiti Teknologi Malaysia and B.IT (Software Engineering) from Universiti Malaysia Terengganu. Her research interests include image enhancement, information security, and software engineering, and she has numerous publications. She won many innovation competitions and was recently awarded the Special Innovation Award in the *Bank Innovation Challenge 2024* for the research grant. She can be reached at sakinah.ali@usim.edu.my and via www.linkedin.com/in/sakinah-ali-pitchay