

КОМПЬЮТЕРНЫЕ СИСТЕМЫ И ИНФОРМАЦИОННЫЕ ТЕХНОЛОГИИ COMPUTER SCIENCE

doi: 10.17586/2226-1494-2022-22-2-254-261

Research on the effectiveness of noise reduction when encoding a lossless speech signal

Tamilselvan Akilan¹✉, Laxmi Raja², Udhayakumar Hariharan³

¹ Galgotias College of Engineering and Technology, Greater Noida, 201310, India

² Karpagam Academy of Higher Education, Coimbatore, 641021, India

³ Chandigarh University, Mohali, 140413, India

¹ t.akilan@galgotiacollege.edu.in, <https://orcid.org/0000-0002-3593-4298>

² laxmirajaphd@gmail.com, <https://orcid.org/0000-0001-6040-8794>

³ hariharan.e11201@cumail.in, <https://orcid.org/0000-0002-3144-2341>

Abstract

In the meantime, speech coding is one of the methods to represent the digital speech signal as in possible fewer bits value and to maintain the quality and its clearness. In omnipresent situations, encryption and examination of speech maintain a crucial role in various acoustic-based coding systems. This paper, using subband and Huffman coding technique, has been used for speech signals description to reduce the occupied by the speech data memory. The amplitude values of the taken speech are segregated after pre-processing, windowing and decomposition techniques. These data are converted into the frequency domain using discrete cosine transform (DCT). Then 90 foremost coefficients have been coded by Huffman method, they contain the most valuable information of speech signals. Signals are segregated then and subband coding techniques applied. To reconstruct the input speech, the taken speech is re-transformed in the form of time-domain applying through inverse discrete cosine transform (IDCT). This experiment is carried out by speech data at 8 kHz with 16 bits/per sample. The SNR (Signal to Noise Ratio) shows the efficiency of this applied technique.

Keywords

decomposition, IDCT, DCT, Huffman, SNR, subband, quantization, windowing

For citation: Akilan T., Raja L., Hariharan U. Research on the effectiveness of noise reduction when encoding a lossless speech signal. *Scientific and Technical Journal of Information Technologies, Mechanics and Optics*, 2022, vol. 22, no. 2, pp. 254–261. doi: 10.17586/2226-1494-2022-22-2-254-261

УДК 004.04

Исследование эффективности шумоподавления при кодировании речевого сигнала без потерь

Тамилсельван Акилан¹✉, Лакшми Раджа², Удхаякумар Харихаран³

¹ Инженерно-технологический колледж Галготиаса, Большая Нойда, 201310, Индия

² Академия высшего образования Карпагама, Коимбатур, 641021, Индия

³ Университет Чандигарха, Мохали, 140413, Индия

¹ t.akilan@galgotiacollege.edu.in, <https://orcid.org/0000-0002-3593-4298>

² laxmirajaphd@gmail.com, <https://orcid.org/0000-0001-6040-8794>

³ hariharan.e11201@cumail.in, <https://orcid.org/0000-0002-3144-2341>

Аннотация

Кодирование речи — один из методов представления цифрового речевого сигнала с использованием малого числа битов, при этом возможно сохранить их качество и точность. В большинстве ситуаций шифрование и качество речи играют решающую роль в различных акустических системах кодирования. Предложен способ уменьшения занимаемой памяти, используемой речевыми данными с применением поддиапазона и алгоритма Хаффмана для речевых сигналов. Выделены значения амплитуды речевого сигнала после предварительной обработки, оконной обработки и применения методов декомпозиции. Полученные данные преобразованы в

© Akilan T., Raja L., Hariharan U., 2022

частотную область с использованием дискретного косинусного преобразования (Discrete Cosine Transform, DCT). Проведено кодирование методами Хаффмана 90 основных коэффициентов, содержащих наибольшее количество информации о речевых сигналах. Для восстановления исходной речи закодированный сигнал повторно преобразован в форму во временной области с применением обратного дискретного косинусного преобразования (Inverse Discrete Cosine Transform, IDCT). Выполнен эксперимент с речевыми данными с 16 битами по выборке на частоте 8 кГц. Величина показателя SNR (отношение сигнал/шум) показывает эффективность предлагаемого метода.

Ключевые слова

декомпозиция, дискретное косинусное преобразование, DCT, обратное дискретное косинусное преобразование, IDCT, алгоритм Хаффмана, SNR, поддиапазон, квантование, оконное преобразование

Ссылка для цитирования: Акилан Т., Раджа Л., Харихаран У. Исследование эффективности шумоподавления при кодировании речевого сигнала без потерь // Научно-технический вестник информационных технологий, механики и оптики. 2022. Т. 22, № 2. С. 254–261 (на англ. яз.) doi: 10.17586/2226-1494-2022-22-2-254-261

Introduction

Coding in speech is the process of the ability to transform the speech signals into a more compressed form with minimum repetition of the speech signal. The speech signal is well transformed into digital form, and it may be stored in a digital medium. It is possible to decode the speech data with the finest achievable good quality [1]. Meanwhile, other various signals, being sampled, have a lot of data that may neither redundant nor perceptually unrelated. The process of splitting the taken speech signals into subbands using bandpass filters, and after that coding every band individually is known as the subband coding technique. The total number of samples to be coded is taken in minimum, and then sampling rate of the speech signals in every band is minimized by the decimation process. Bandpass filters aren't perfect because of a few overlaps in the midst of the adjacent bands, and errors may happen during the decimation process. Subband coding plays a vital role so that each individual band can be coded differently. It was the main benefit of subband coding used in this paper, and this method may control the error in coding for each band to manage properly the human hearable way.

In Huffman coding, the speech signals are arranged in the decreasing order of frequency (increasing order of probability of occurrence) [2]. The salient components of a Huffman tree coding are nodes as well as leaves. At every stage, the calculation has been done for the two leaves which are having the lowest probability value. After that, it may couple together to create a new node. In this process, the tree has been constructed in the way of the bottom-up approach on $N-1$ steps. Let N will be the number of symbols. Value 0 is allocated to every left going path, and each right going path is assigned 1.

Proposed Methodology

Speech Compression Technique

This technique may differ in the amount of compression in data, as well as in the sample rate used. It provides distinguishing levels of system complexity and minimizes the quality of speech information.

Then, the stored compressed waveform will be transferred along with loss or without loss. Speech signals are managed through mingling, equalization, and filtering. The audio signal enters the encoder that use only lesser bits rather than original speech signals bit-rate. As a result, the

transmission bandwidth of the speech signal is minimized, as well as minimized the memory size occupied by the speech files. Basically, the speech compression technique is divided into two different categories: lossless and lossy coding. Lossy coding is truly clear to hear for human perceptibility, and lossless compression has a factor of compressing ranges from 6 to 1. Fig. 1 shows the original speech signal. Fig. 2, *a* describes the block diagram of the proposed speech coding by the Huffman technique. In that technique the input speech signal is decomposed into eight levels of segmentation; then windowing is applied while a lengthy signal was increased by the windowing function with determinate length, providing very little weighted length in the form of input speech data without coding. Then it will be reconstructed by the discrete cosine transform method to find the small-sized frames and will be assembled in the matrix form. The Discrete Cosine Transform (DCT) process can be applied to the produced matrix. The elements are arranged in the appropriate matrix format to identify the components and index values. Herewith, a total of 90 values of speech data had been used and entered into four processing units to obtain the good hearing speech signal. The values are arranged in higher to lower order, then the bigger values and also threshold values have been taken for further processing. Those threshold values are quantized in order to convert the sample analogue signal into voltage value, and further to convert into a binary digit, which will be read by the computer system. Finally, Huffman technique was applied to the binary digit to obtain the compressed speech signals. Afterwards, Inverse Discrete Cosine Transform (IDCT) will be applied to get the decompressed speech signals.

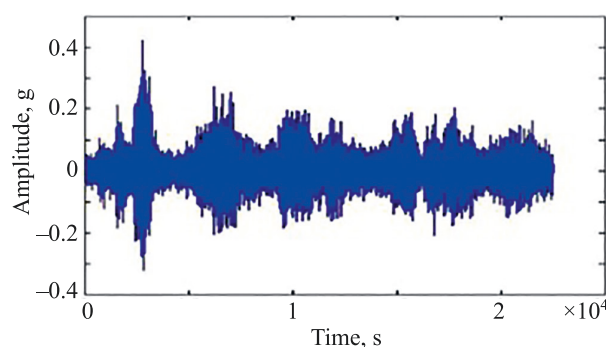


Fig. 1. Original speech signals

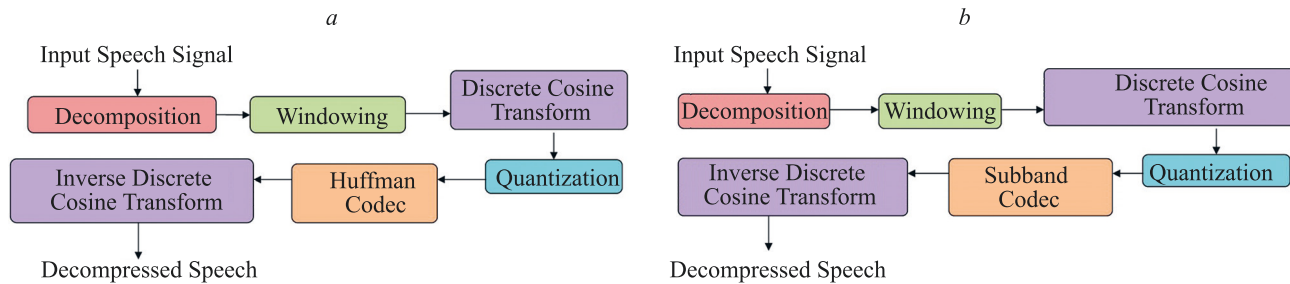


Fig. 2. Block diagrams of proposed speech coding by Huffman technique (a) and of speech coding using subband technique (b)

Fig. 2, b describes the same process which was taken in Huffman technique here, but instead of Huffman coding, the subband coding has been done to get the compressed speech signals. Finally, both techniques were compared to view what scenario and which technique will produce the better results.

Huffman Coding Technique

With variable-length codes, Huffman coding is a popular approach for data compressing. The approach produces a set of variable-length codewords with the smallest average length and assigns them to the symbols producing a collection of data symbols (an alphabet) with their frequencies of occurrence (or, equivalently, their probabilities). Huffman procedure builds a code tree from the ground up (and the bits of each codeword are constructed from right to left). The sampling rate for the signals in each band is decimated to restrict the number of samples to be coded to a minimum. The basic idea is that a discrete unitary transform is applied to a set of speech samples, and the resultant transform coefficients are quantized and coded for transmission to the destination. Because more bits may be allocated to the perceptually essential coefficients, low bit rates and high performance can be achieved.

This Huffman coding provides the source for more research areas. This coding technique develops a code tree with a bottom-up level, and each bit of codeword are defined from right to left side [3]. To optimize the number of data samples to be encoded, the sampling rate of the signals in every band is minimized by the decimation technique. The salient method is that a less part of speech

samples is going to be processed through discrete unitary transform and, in result, the transform coefficients were quantized, and then it can be encoded for further transmission of the beneficiary. Lower bit rates and high bit rates can be found due to a greater number of bits is to be assigned to the omniscience required coefficients. Fig. 3 shows the example of Huffman coding technique.

- Expected size
 - Original $\Rightarrow 13 \times 2 + 25 \times 2 + 50 \times 2 + 12 \times 2 = 2.00$ bits/symbol
 - Huffman $\Rightarrow 13 \times 3 + 10 \times 2 + 0 \times 1 + 12 \times 3 = 1.75$ bits/symbol

Above, an example of calculating compressed size is shown which declares the original speech signals having the highest memory of 2 bits per symbol while the Huffman encoding technique will produce the compressed memory size is 1.75 bits per symbol. This means that the Huffman encoding will give good result compared to the original encoding process. An example has been clearly calculated, and the same has been shown in the Table and in Fig. 3. The symbols are organized in decreasing frequency order (increasing order of probability of occurrence) [4]. Nodes and leaves are the most important parts of a Huffman tree. We calculate the two leaves with the lowest probability at each step and then group them together to construct a node. In this way, the tree is built from the bottom up in $N-1$ steps, where N is the number of symbols. Here a1, a2, a3, a4 were taken as an example, for which the original speech compression has been implemented; the data in fact is available in terms of the results and discussion. Symbol a1 (can be any character) is assigned to each left-going path, whilst symbol a2 is assigned to each right-going path. Symbol a2 (can be any character) is assigned to each right heading path in constructing the assigned code. In order to construct the code corresponding to a given symbol, move down the tree in a top-down approach and build up the code for that symbol.

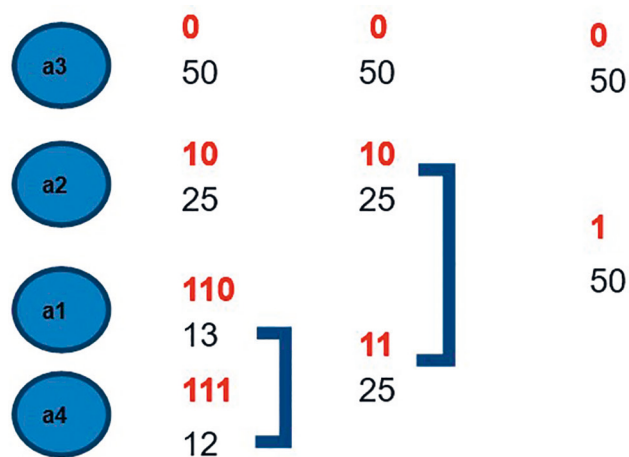


Fig. 3. Example of Huffman coding Technique

Table. Comparison of Huffman coding with the original speech coding

Symbol	a1	a2	a3	a4
Frequency	13	25	50	12
Original encoding	00	01	10	11
	2 bits	2 bits	2 bits	2 bits
Huffman encoding	110	10	0	111
	3 bits	2 bits	1 bit	3 bits

Subband Coding Technique

The subband coding technique is the process of splitting the original speech data into sub-signals. This can be done by utilizing the various bandpass filters applying then the speech coding technique to each and every signal separately. The whole process is known as the subband coding technique [5]. While maintaining the more samples that are going to be coded with the very few data, the sampling rate of each signal in every band is minimized through the decimation method. However, the bandpass filters weren't perfect because a few redundancies or overlap between adjacent bands may occur, and aliasing may occur in the process of the decimation method. Fig. 4 shows the processes of subband encoding (Fig. 4, *a*) and decoding (Fig. 4, *b*). On the picture, $x(n)$ is input speech signals, $y(n)$ is output speech signals, $h(z)$ is impulse response and n is the level. Filters H_0, H_1, G_0, G_1 are the low decomposition, high decomposition, low reconstruction and high reconstruction procedures. The synthesis filter G is a time reversal version of the analysis filter H . Symbol $\downarrow 2$ (Fig. 4, *a*) is up-sampling and symbol $\uparrow 2$ (Fig. 4, *b*) is down-sampling. Fig. 4, *a* shows the subband encoding process describing that the input signals $x(n)$ passed the pair of high-pass filter and down-sampling in order to obtain the output $y(n)$. Fig. 4, *b* shows the subband decoding process in which the output $y(n)$ will be taken and processed with the pair of up-sampling and high pass filter in order to obtain the original speech signal. The output signals $y_0(n), y_1(n), y_2(n), y_3(n)$ have been fed as an input of subband encoding technique in order to decode the speech signal to get the decompression speech signal $\hat{x}(n)$.

Process of Decomposition

Wavelets will be decomposed for a speech signal into different resolutions or into different frequency

bands. Speech compression was considered in terms of choosing a smaller amount of approximation coefficients, and then a few detailed coefficients can more accurately define the signal components [6]. Initial speech signals are primarily examined to concentrate on the spoken and unspoken parts of the speech signal. The speech signal equivalent to different noisy areas is used as the first 8 level approximations. The 8 level decomposition of clean speech signals is shown in Fig. 5, *a-h*.

Windowing

The windowing technique (Fig. 6) is applied to untreated speech frames to facilitate minimizing spectral leakage. A larger partition of speech signal processing is carried out in this way by taking short windows (overlapping may occur). A frame is commonly called as the short window of the signal. A lengthy signal was increased by the windowing function with determinate length, providing very little weighted length in the form of input speech data without coding.

DCT/IDCT

Input speech signal of the vector data is separated into small-sized frames and assembled in the form of matrix format. DCT process can be applied to the produced matrix. The elements are arranged in the appropriate matrix format to identify the components and index values [7, 8]. Herewith, a total of 90 values of speech data had been used and entered into four processing stages to obtain the good hearing speech signal. The values are arranged in higher to lower order. For further processing, the bigger values and also threshold values are taken. As a result, some of the values are discarded due to coefficients that are below the threshold values. Consequently, a reducing the size of the speech signal will mean that the compression has been applied to the speech signal.

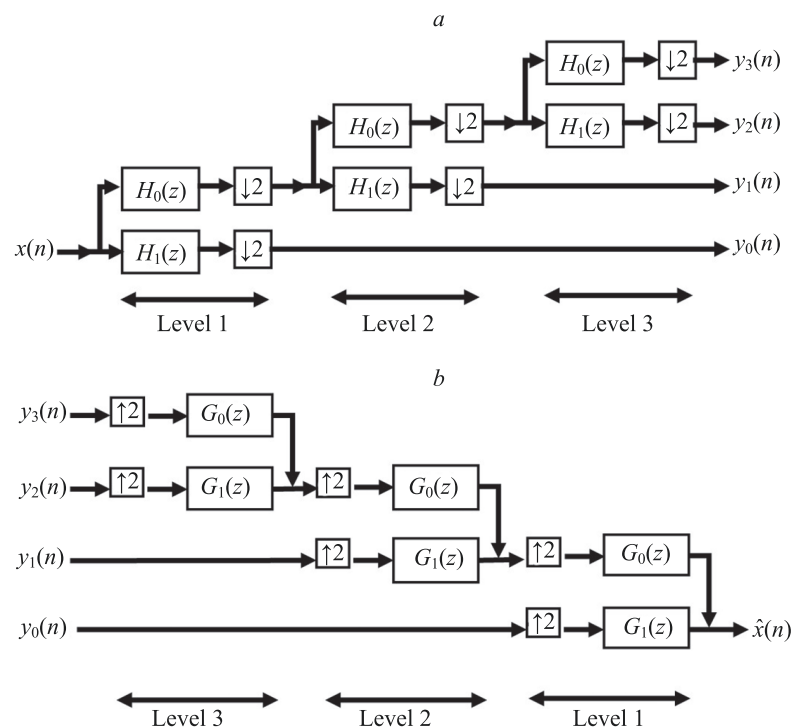


Fig. 4. The processes of subband encoding (*a*) and decoding (*b*)

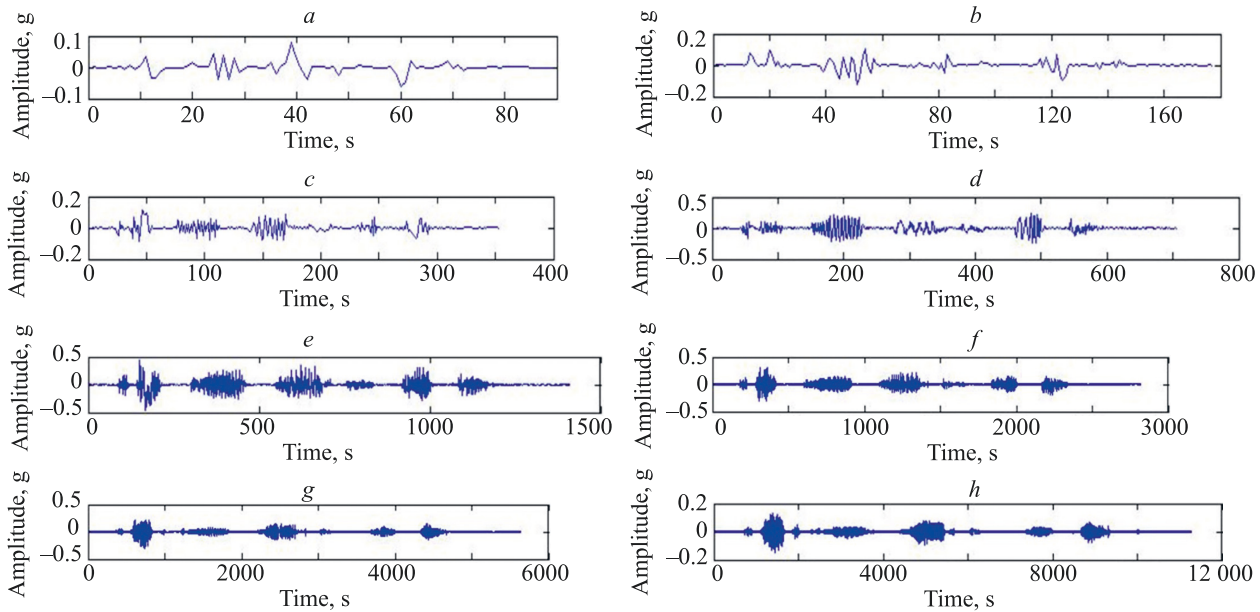


Fig. 5. Eight level decomposition of clean speech signals: level 1 (a); level 2 (b); level 3 (c); level 4 (d); level 5 (e); level 6 (f); level 7 (g); level 8 (h)

$$X_c(k) = \left(\frac{1}{N} \right) \sum_{n=0}^{N-1} x_n \cos\left(\frac{k2\pi n}{N}\right),$$

where $k = 0, 1, 2, N-1$, k is the length of the filter.

The backward recovery of the original speech data will be done by applying the reconstruction technique, i.e. IDCT will be applied to the speech signal. By this process, the speech signal will be reconstructed.

$$X_c(k) = \sum_{n=0}^{n-1} c[u] X_n \cos\left(\frac{k2\pi n}{N}\right),$$

where X_n is the result of DCT; C is discrete frequency variables ($0, 1, 2, N-1$); $k = 0, 1, 2, 3$, $n-1$; u is transform-domain horizontal frequency coordinates. Thuswise, $c[u] = 1$ if $u = 0$, and $c[u] = 2$ if $u = 1, 2, 3, N-1$. This predefined form of DCT $[u]$ is denoted for iterative horizontal coordinates.

Quantization Technique

The sample analogue signal will be transformed through the voltage value into a binary digit, which will be read by

the computer system [9, 10]. The convention from infinitely precise amplitude into the binary digit is known as the quantization technique.

Discussion of Results

The speech signals are taken as noisy, obtained from the different locations like a car, exhibition, airport, babble, restaurant, railway station, street, and train station. Those noises are processed and compressed using DCT coefficients along with subband coding and Huffman coding technique. The speech signal with noise is shown in Fig. 7.

The clean input signal has the original data of the user speech in various locations. Fig. 8 shows DCT applied to the noisy speech signal.

By DCT processing, the data is arranged through the matrix form in order to identify the components as well as the indices. This processing has been completed and then only the threshold value will be fixed. The coefficient value which is below the threshold will be discarded. Compression technology will be applied in order to reduce the signal size.

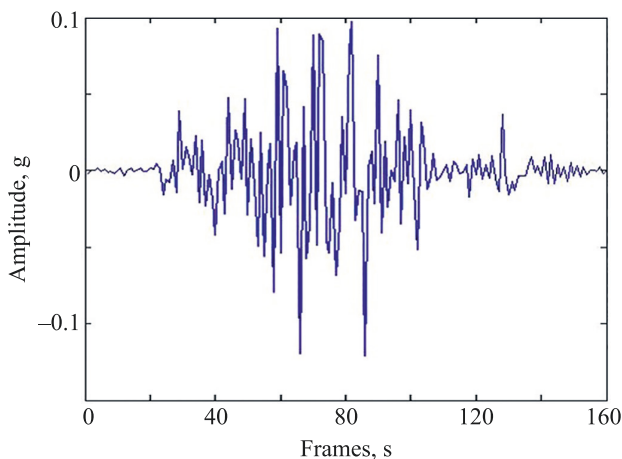


Fig. 6. Hamming window

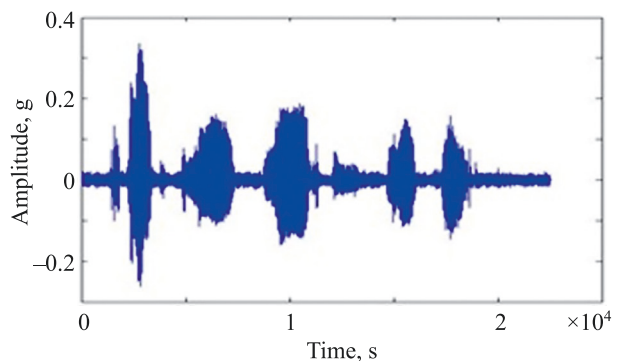


Fig. 7. Speech signal with noise

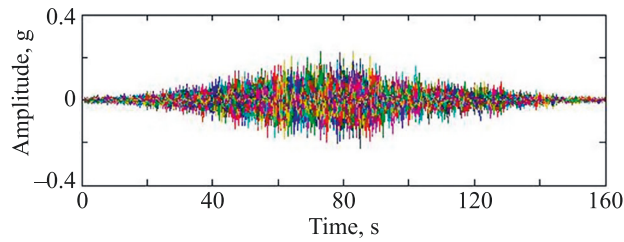


Fig. 8. DCT applied to the noisy speech signal

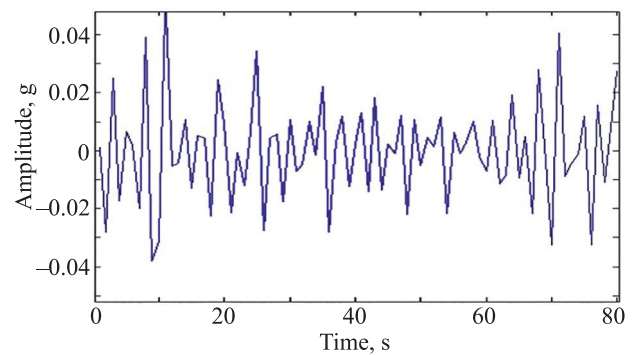


Fig. 9. Reconstructed noisy signal

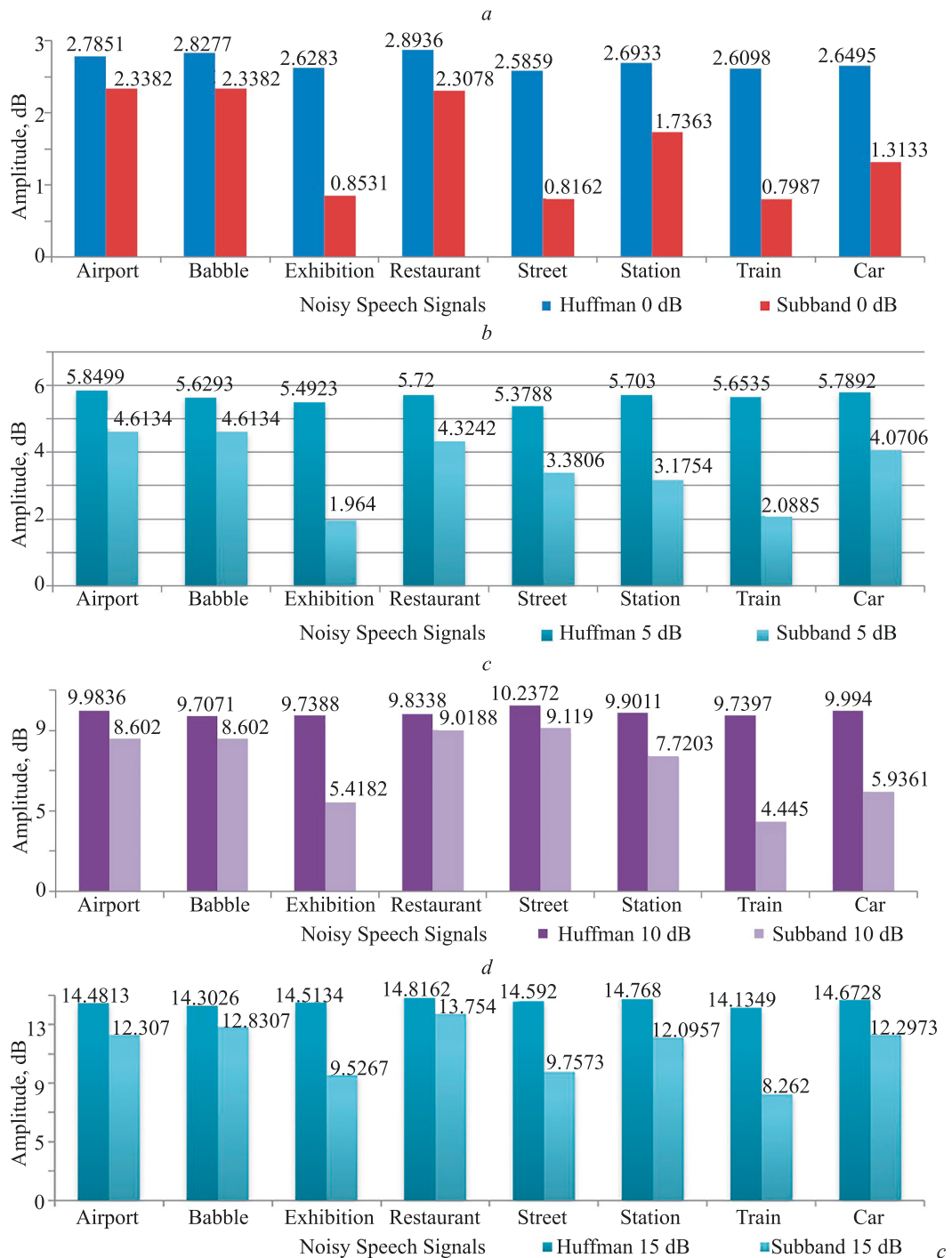


Fig. 10. SNR comparison of Huffman with Subband with various decibels: 0 dB (a), 5 dB (b), 10 dB (c), 15 dB (d)

By the reconstruction technique, the values will be transformed back into their original form via the threshold values (Fig. 9). The original frequency is almost nearer to the taken input speech signal with an accuracy of 85 %.

Database Details

These processes are applied to the speech signal with noisy speech, and the used database is **NOIZEUS**. This database may contain noisy signals. Those signals are taken in various locations like restaurants, train stations, airports, babble, exhibitions, streets, cars. Then additional noise is added to the input speech signals, and then this mixture will be processed. At last, the Signal to Noise Ratio (SNR) value has been used, and the result will be compared to the clear speech [8]. The SNR Comparison of Huffman with Subband with various decibels (dB) is shown in the Fig. 10.

Conclusion

Coding in speech is a recent finding research area, and compression of the speech signal is a standard way for making and reducing the speech signal in the compression form. This paper mainly looks into developing an effective speech coding technique using Huffman and Subband

coding. DCT/IDCT based speech compression approaches are used to get better results. This kind of process that gives a result is carried out with the database named NOIZEUS. Input speech is taken into the reconstruction process from the encoded features. The subband coding and Huffman technique worked efficiently, and it provides a better result in these taken complicated situations. For some speech signals, it was easy to identify each and every word even in the distorted utterance. During listening, the compressed speech signal is obtained with good audibility quality. While using these two taken techniques, the original frequency is obtained almost nearer to the taken input speech signal with an accuracy of 85 %. The improvement result is obtained, whilst applying subband coding for 0 dB speech signal to the noise speech signal like exhibition, street and train. The improvement result is obtained whilst applying Huffman Coding for 5 dB speech signal to the noisy speech signal like babble, exhibition, street and train. The improvement result is obtained whilst applying Huffman Coding for 10 dB speech signal to the noisy speech signal like airport, restaurant, street, train and car. The improvement result is obtained whilst applying Huffman Coding for 15 dB speech signal to the noisy speech signal like airport, babble, station and car.

References

1. Lv S., Hu Y., Zhang S., Xie L. DCCRN+: channel-wise subband DCCRN with SNR estimation for speech enhancement. *Proc. of the 22nd Annual Conference of the International Speech Communication Association (INTERSPEECH)*, 2021, pp. 2816–2820. <https://doi.org/10.21437/Interspeech.2021-1482>
2. Taujuddin N.S.A.M., Ibrahim R., Sari S. Image compression using a new adaptive standard deviation thresholding estimation at the wavelet details subbands. *Proc. of the 2nd International Conference on Computing Technology and Information Management (ICCTIM)*, 2015, pp. 109–114. <https://doi.org/10.1109/ICCTIM.2015.7224602>
3. Pal R. Speech compression with wavelet transform and Huffman coding. *Proc. of the 4th International Conference on Communication, Information and Computing Technology (ICCICT)*, 2021, pp. 1–4. <https://doi.org/10.1109/ICCICT50803.2021.9510116>
4. Li S., Zheng Z., Dai W., Xiong H. Lossy image compression with filter bank based convolutional networks. *Proc. of the Data Compression Conference (DCC)*, 2019, pp. 23–32. <https://doi.org/10.1109/DCC.2019.00010>
5. Cooper C., Marcellin M. Lossless wideband RF compression via lifting-based IIR subband decomposition. *IEEE Transactions on Aerospace and Electronic Systems*, 2020, vol. 56, no. 1, pp. 823–829. <https://doi.org/10.1109/TAES.2019.2919436>
6. Vatsa S., Dr. Sahu O.P. Speech compression using discrete wavelet transform and discrete cosine transform. *International Journal of Engineering Research & Technology (IJERT)*, 2012, vol. 1, no. 5, pp. 1–6.
7. Balaji V.R., Subramanian S. A novel speech enhancement approach based on modified DCT and improved pitch synchronous analysis. *American Journal of Applied Sciences*, 2014, vol. 11, no. 1, pp. 24–37. <https://doi.org/10.3844/ajassp.2014.24.37>
8. Vats S., Rathee G. An image-compression decomposition analysis of sub-bands using threshold implementation. *Proc. of the 3rd International Conference on Image Information Processing (ICIIP)*, 2015, pp. 366–369. <https://doi.org/10.1109/ICIIP.2015.7414797>
9. Luneau J.-M., Lebrun J., Jensen S.H. Complex wavelet modulation subbands for speech compression. *Proc. of the Data Compression Conference (DCC)*, 2009, pp. 457. <https://doi.org/10.1109/DCC.2009.52>
10. Mack W., Habets E.A.P. Deep filtering: Signal extraction and reconstruction using complex time-frequency filters. *IEEE Signal Processing Letters*, 2020, vol. 27, pp. 61–65. <https://doi.org/10.1109/LSP.2019.2955818>

Литература

1. Lv S., Hu Y., Zhang S., Xie L. DCCRN+: channel-wise subband DCCRN with SNR estimation for speech enhancement // *Proc. of the 22nd Annual Conference of the International Speech Communication Association (INTERSPEECH)*. 2021. P. 2816–2820. <https://doi.org/10.21437/Interspeech.2021-1482>
2. Taujuddin N.S.A.M., Ibrahim R., Sari S. Image compression using a new adaptive standard deviation thresholding estimation at the wavelet details subbands // *Proc. of the 2nd International Conference on Computing Technology and Information Management (ICCTIM)*. 2015. P. 109–114. <https://doi.org/10.1109/ICCTIM.2015.7224602>
3. Pal R. Speech compression with wavelet transform and Huffman coding // *Proc. of the 4th International Conference on Communication, Information and Computing Technology (ICCICT)*. 2021. P. 1–4. <https://doi.org/10.1109/ICCICT50803.2021.9510116>
4. Li S., Zheng Z., Dai W., Xiong H. Lossy image compression with filter bank based convolutional networks // *Proc. of the Data Compression Conference (DCC)*. 2019. P. 23–32. <https://doi.org/10.1109/DCC.2019.00010>
5. Cooper C., Marcellin M. Lossless wideband RF compression via lifting-based IIR subband decomposition // *IEEE Transactions on Aerospace and Electronic Systems*. 2020. V. 56. N 1. P. 823–829. <https://doi.org/10.1109/TAES.2019.2919436>
6. Vatsa S., Dr. Sahu O.P. Speech compression using discrete wavelet transform and discrete cosine transform // *International Journal of Engineering Research & Technology (IJERT)*. 2012. V. 1. N 5. P. 1–6.
7. Balaji V.R., Subramanian S. A novel speech enhancement approach based on modified DCT and improved pitch synchronous analysis // *American Journal of Applied Sciences*. 2014. V. 11. N 1. P. 24–37. <https://doi.org/10.3844/ajassp.2014.24.37>
8. Vats S., Rathee G. An image-compression decomposition analysis of sub-bands using threshold implementation // *Proc. of the 3rd International Conference on Image Information Processing (ICIIP)*. 2015. P. 366–369. <https://doi.org/10.1109/ICIIP.2015.7414797>
9. Luneau J.-M., Lebrun J., Jensen S.H. Complex wavelet modulation subbands for speech compression // *Proc. of the Data Compression Conference (DCC)*. 2009. P. 457. <https://doi.org/10.1109/DCC.2009.52>
10. Mack W., Habets E.A.P. Deep filtering: Signal extraction and reconstruction using complex time-frequency filters // *IEEE Signal Processing Letters*. 2020. V. 27. P. 61–65. <https://doi.org/10.1109/LSP.2019.2955818>

Authors

Tamilselvan Akilan — M.E, Assistant Professor, Galgotias College of Engineering and Technology, Greater Noida, 201310, India, [sc](https://orcid.org/0000-0002-3593-4298) 56801096100, <https://orcid.org/0000-0002-3593-4298>, t.akilan@galgotiacollege.edu

Laxmi Raja — PhD, Assistant Professor, Karpagam Academy of Higher Education, Coimbatore, 641021, India, [sc](https://orcid.org/0000-0001-6040-8794) 57197747072, <https://orcid.org/0000-0001-6040-8794>, laxmirajaphd@gmail.com

Udhayakumar Hariharan — PhD, Assistant Professor, Chandigarh University, Mohali, 140413, India, [sc](https://orcid.org/0000-0002-3144-2341) 57216226566, <https://orcid.org/0000-0002-3144-2341>, hariharan.e11201@cumail.in

Авторы

Акилан Тамилсельван — М.Е., доцент, доцент, Инженерно-технологический колледж Галготиаса, Большая Нойда, 201310, Индия, [sc](https://orcid.org/0000-0002-3593-4298) 56801096100, <https://orcid.org/0000-0002-3593-4298>, t.akilan@galgotiacollege.edu

Раджа Лакшми — PhD, доцент, доцент, Академия высшего образования Карпагама, Коимбатур, 641021, Индия, [sc](https://orcid.org/0000-0001-6040-8794) 57197747072, <https://orcid.org/0000-0001-6040-8794>, laxmirajaphd@gmail.com

Харихаран Удхаякумар — PhD, доцент, доцент, Университет Чандигарха, Мохали, 140413, Индия, [sc](https://orcid.org/0000-0002-3144-2341) 57216226566, <https://orcid.org/0000-0002-3144-2341>, hariharan.e11201@cumail.in

Received 04.09.2021

Approved after reviewing 04.02.2022

Accepted 17.03.2022

Статья поступила в редакцию 04.09.2021

Одобрена после рецензирования 04.02.2022

Принята к печати 17.03.2022



Работа доступна по лицензии
Creative Commons
«Attribution-NonCommercial»