

A Novel Method for Separation of Two Speech Signals Recorded with Two Microphones

Mahbanou Zohrevandi

Department of Computer
Engineering, Science and
Research Branch, Islamic Azad
University, Tehran, Iran.
m.zohrevandi@srbiau.ac.ir

Saeed Setayeshi*

Department of Medical
Radiation Engineering, Faculty
of Energy Eng and Physics,
Amirkabir University of
Technology, Tehran, Iran.
setayeshi@aut.ac.ir

Azam Rabiee

Department of Computer
Science, Dolatabad Branch,
Islamic Azad University,
Isfahan, Iran.
a.rabiee@iauda.ac.ir

Midia Reshadi

Department of Computer
Engineering, Science and
Research Branch, Islamic Azad
University, Tehran, Iran.
reshadi@srbiau.ac.ir

Received: 21 September 2020 - Accepted: 29 November 2020

Abstract—The objective of this study was to develop a technique for differentiation of two instantaneous audio signals recorded concurrently by two microphones. To achieve this objective, a signal was first subjected to Short Time Fourier Transform (STFT), then each frequency bin was processed individually by a two-phase technique consisting of an estimation of orthogonalization matrix and rotation matrix in that order. The proposed method introduces a new geometric approach to the separation of sparse signals. In fact, it introduces two new steps for the general BSS technique. In fact, these two steps are presented in this paper instead of the whitening and rotating steps in blind speech separation algorithms. This paper separates audio sources in difficult situations with significant improvements in algorithm performance. Therefore, the experiments were performed in difficult conditions, such as when the microphones are collinear, short distances of microphones and having short sources. Therefore, the experiments were performed in difficult conditions. Performance of the proposed method was evaluated by a number of numerical examples. The experiments were performed using the Roosim simulator, and the results showed that the proposed algorithm is a simple and useful solution for separating two speech signals recorded with two microphones.

Keywords—Blind source separation; sparse signals; orthogonalization; rotation.

I. INTRODUCTION

In the past few decades, the cocktail party problem, which refers to differentiation of multiple speech signals from recordings of noisy settings, have found potential applications in hearing aids, speech detection, and teleconferencing, and have therefore received growing attention from the research community [1]. One method with

significant potential in this field is Blind Source Separation (BSS), which refers to the extraction of several statistically independent audio signals from a recording in the absence of any data concerning the mixing conditions [1-4]. The potential applications of BSS include, but are not limited to, speech processing, wireless communication, medical analysis, and remote sensing [5]. The methods proposed for BSS problem can be categorized into

* Corresponding Author

two groups depending on the assumed propagation model: instantaneous mixture [6-9] and convolutive mixture [5,9,10]. Among the major statistical solutions, and perhaps the most prominent, developed for the BSS problem is Independent Component Analysis (ICA). The mixed signals of an instantaneously mixed recording can be decomposed by instantaneous ICA [10].

Convolutive mixing refers to the mixing of signals in the spaces where sound propagation and its reflection of different objects cause some time delays in signals. In reverberation spaces, sound reflecting off walls and objects can cause a time dependency in the recorded signal mixture. Because of this time dependency, the BSS problem must be solved through various methods, which can be divided into two main categories based on their approach to ICA utilization [11]. The first category is known as time-domain BSS. In this category, the model of convolutive mixtures is directly related to ICA. However, the long FIR filters required for the convolution operation makes the convolutive ICA more complex, and computation intensive, than its instantaneous counterpart [12].

The second category is known as frequency-domain BSS, which allows the ICA algorithm to be implemented on a simple frequency domain and executed for each frequency in an independent manner [13]. This approach however brings up the problem of permutation ambiguity in the ICA solution. To address this issue, the permutation in each frequency bin should be adjusted so that the signal extracted in the time reverberation settings, the audio signals, can be predictively convolutive [4]. Meanwhile, the amplitude modulations in the voiced parts and the intermixing of voiced and unvoiced parts of phrases make the human speech a highly nonstationary signal. The common practice for converting a time-domain signal into its frequency-domain counterpart is to use the sliding window discrete Fourier transform (DFT) or the short-time Fourier transform (STFT). When following this approach, it is imperative to select the window length in a way that the data of each window would remain quasi-stationary. Speech separation is also associated with four other properties of speech signals: i) In an acoustic environment, speech signals of different speakers at different locations are statistically independent; ii) Temporal structure of each given speech signal is normally unique over short time frames (< 1 second); iii) In small time frames (25 ms) speech signals are quasi-stationary, but in longer frames they become completely nonstationary [4, 14-17].

The BSS problem is an extensively studied subject of research, and numerous solution algorithms have been proposed for variants of this problem. The authors of [15] have developed an FIR backward model with adequate ability to separate

distinct model sources. In [16], a two-step approach has been utilized for Monaural Speech Separation; in the first stage harmonicity has been utilized to separate the auditory spectrogram of the mixture into its components and, in the next stage, a top-down process based on harmonic filters has been used to enhance the quality of the separated signals. In this study, first a forward model has been estimated by a least squares optimization technique, with a multi-path channel identification then carried out, accordingly. The frequency-domain technique proposed in [17] utilizes PARAllel FACtor (PARAFAC) analysis for multichannel BSS of convolutive speech mixtures. In this method, computational complexity is controlled by incorporating a dimensionality reduction step into the PARAFAC algorithm.

In the present paper, speech separation problem is explored in a scenario where there is a presence of two microphones and two speakers. Fig. 1 shows the schematic representation of the proposed speech separation algorithm. As can be seen, this algorithm first segments the signal mixture, as the non-stationary nature of speech signals limits the validity of statistical analyses to short quasi-stationary frames. As a result, signal mixtures are first segmented into short time frames (25 ms), and then the signals within the frequency domain ($X(w; q)$) of each frame are extracted. The proposed signal separation approach consists of two phases of orthogonalization and rotation. Also, the scaling and permutation ambiguities are addressed once all frequencies are subjected to BSS process. This algorithm has no information about resources and therefore is not compared to supervised learning methods and is compared to algorithms that use blind data. In [7], with a simple algorithm, it detects time-frequency points with only one active source in instantaneous mixtures and then estimates a mask for each source. [25] In the frequency domain, using independent component analysis (ICA), it estimates the mixing filters and then solves the permutation problem with a Recursively Regularized implementation of the ICA (RR-ICA).

New research on speech separation is known as supervised learning methods and includes learning patterns of speech, speakers, and background noise [18]. But since we do not have any information to train learning algorithms on blind problems, these algorithms cannot be used to solve blind speech separation problems. Also training these learning algorithms involves a large amount of data processing. The advantages of the proposed technique over its rivals include simplicity and ease of use, high speed of computation, and the absence of convergence failure issues.

The contributions of this algorithm are as follows. First, it separates instantaneous audio signals with a new geometric approach. This algorithm consists of

two steps. These two steps are presented in this paper instead of the bleaching and rotation steps in blinded speech separation algorithms. Second, this paper uses the rotation method to separate audio sources instead of estimating the mixing matrix. Third, in difficult separation situations, such as when the microphones are collinear, the microphones are close together and short sources work well. Fourth, this algorithm can separate more than two speakers, but to separate both speakers, the algorithm must be run once and the efficiency of the algorithm will decrease.

In the remainder of this paper, Section 2 provides a general formulation for convolutive BSS problem; Section 3 describes the proposed technique; Section 4 provides the results of numerical simulations; and Section 5 concludes the paper.

II. BSS OF CONVOLUTIVE MIXTURES

Assume N statistically independent voice signals denoted by $s(t) = [s_1(t); s_2(t); \dots; s_N(t)]$, which have been recorded by the same number of microphones (N) to form the mixture $x(t) = [x_1(t); x_2(t); \dots; x_N(t)]$. With this assumption, the formulation of noise-free convolutive model in the time domain will be as follows [1,19]:

$$x(t) = H * s(t) = \sum_{l=1}^{L-1} H(l)s(t-l) \quad (1)$$

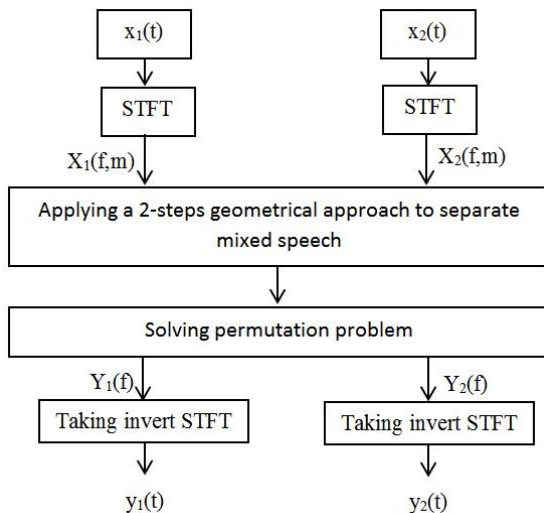


Figure 1. The block diagram of the algorithm.

In this relationship, $*$ is the operator of linear convolution. The $N \times N$ matrix $H(l)$ denoting the mixing system at time-lag l consists of elements h_{nm} denoting the coefficients of the space impulse response for source n and microphone m , which has been modelled as a FIR filter. L is the maximum length of the channel. With these definitions, source signals can be approximated by finding the approximate $N \times N$ inverse-channel matrix W for which:

$$s(t) = W(t) * s(t) = \sum_{k=0}^{K-1} W(k)s(t-k) \quad (2)$$

In the above relationship, K denotes the length of inverse-channel impulse response. This problem can be solved by the use of a time-domain method or a frequency-domain approach. In case of using a time-domain method, K must be greater than the unidentified true channel in order to account for all reflections, and much greater than L to allow for precise evaluation. When using a time-domain approach, one must pay particular attention to the sensitivity of these methods, to channel-order mismatch, and the lack of sufficiently understood identifiability properties, particularly in the presence of underdetermination.

A convolutive mixture in the time domain can be assumed equivalent to instantaneous mixtures in the frequency domain, so the BSS problem of reverberation settings can be solved by using an ICA algorithm in the frequency domain. After subjecting (1) to STFT, the model is obtained from the following relationship [20]:

$$X(w, q) = H(w)s(w, q) \quad (3)$$

In the above relationship, w denotes the angular frequency, and q denotes the frame index. The differentiation process in each frequency bin is given by:

$$Y(w, q) = W(w)X(w, q) \quad (4)$$

In the above relationship $S(w, q) = [S_1(w, q); \dots; S_N(w, q)]$ denotes the source signal within the frequency bin w , is the observed signal, denotes the approximated source signal, and $W(w)$ is the separation matrix for the frequency bin w . Here, the aim is to obtain the $W(w)$ that leads to $Y_i(w, q)$.

III. THE PROPOSED BSS ALGORITHM

This section presents the proposed BSS algorithm for two sparse signals recorded in reverberation settings. Assume the following system of linear equations,

$$X_f(w) = HS_f(w) \quad (5)$$

where $H = [h_{ij}]_{2 \times 2}$ is the mixing matrix, $S_f(w) = [S_1(w); S_2(w)]^T$ is the source vector, and $X_f(w) = [X_1(w); X_2(w)]^T$ is the observation vector. Here, observation vector comprises the STFT domain of two speech signals.

For the BSS problem to be solved and speech signals to be differentiated, separation matrix W needs to be such that elements of the output vector $(Y_f(w) = WX_f(w))$ become as similar as possible to elements of the source vector $S(w)$. For this to become possible, elements of the source vector are assumed to be sparse signals.

Since the source vector S_f comprises the STFT domain of two speech signals s_1 and s_2 , Fig. 2(a) can

be regarded as an accurate representation of the joint distribution of the source vector in the $(S_1; S_2)$ plane. Fig. 2(b) shows the joint distribution of the observation vector when X_f is acquired by the matrix transformation of vector S_f .

The method proposed in this paper follows a geometrical approach to differentiation of sparse signals, which is derived from the general BSS technique shown in Fig. 3 [20]. The proposed method consists of two phases. In the first phase observation will be orthogonalized, and in the second phase the orthogonalized data will be rotated.

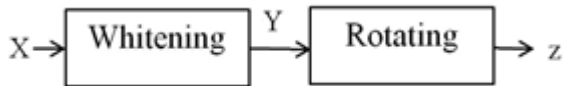


Figure 2. Block diagram of a common BSS method [20].

Whitening of random data mirrors closely the concept of orthogonalization of a set of vectors. Orthogonalizing two vectors involves mapping them observations (c) orthogonalized data to a new set of vectors using a linear transformation so that the inner product between them is zero [25]. Furthermore, rotating data mirrors closely the concept of the shifting of angles in a set of vectors, with respect to special coordinate axes. The orthogonalization part in the proposed algorithm involves finding a new basis in which two vectors illustrated in Fig. 3(b) are perpendicular. The proposed algorithm uses the Gram-Schmidt algorithm [21] to find an orthogonal basis. The Gram-Schmidt algorithm for two vectors x_1 and x_2 is as follows:

$$\begin{cases} o_1 = x_1 \\ o_2 = x_2 - (o_1^T x_2 / o_1^T o_1) o_1 \end{cases} \quad (6)$$

where o_1 and o_2 are two vectors that are perpendicular to each other. After finding a matrix transformation (M) to make vectors x_1 and x_2 perpendicular, the orthogonalized data vector (o) is given by:

$$O = Mx \quad (7)$$

To separate speaker signals from orthogonalized data, we have to rotate the two vectors shown in Fig. 4 in such a way that they are placed on the coordinate axes. A geometrical rotation is a transformation under which the lengths of the original vectors and angles between them do not change. To achieve this goal, we consider the following rotation matrix $R = [r_{ij}]_{2 \times 2}$:

$$R = \begin{bmatrix} \cos \theta & \sin \theta \\ \sin \theta & -\cos \theta \end{bmatrix} \quad (8)$$

Where θ is the angle between vector o_1 and y-axis in the Cartesian coordinate. Finally, the separating matrix is given by:

$$H = RM \quad (9)$$

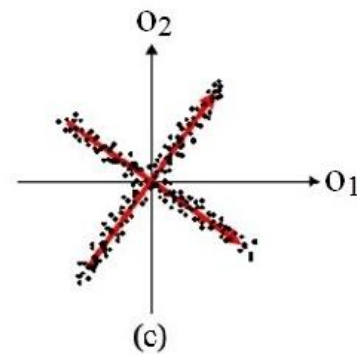
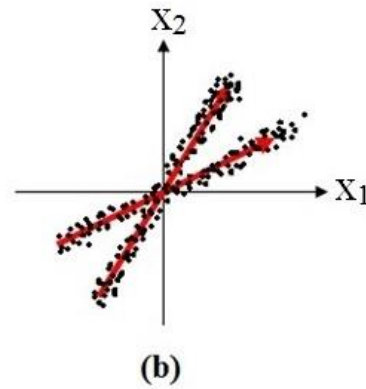
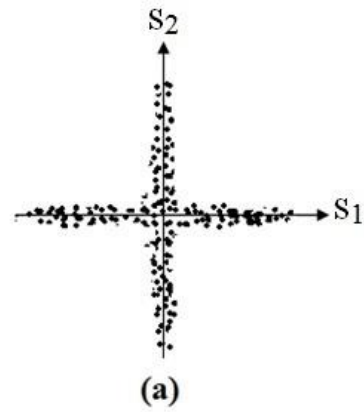


Figure 3. Scatter plot of (a) speaker signals (b).

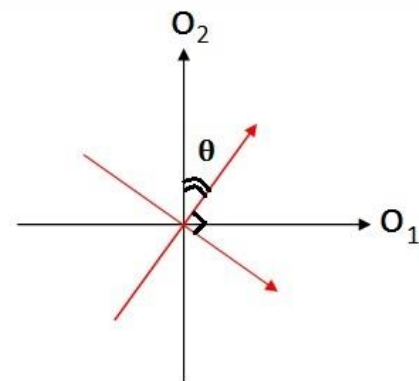


Figure 4. Orthogonalized vectors.

The two-step geometrical approach of the proposed method shown in Fig. 5.

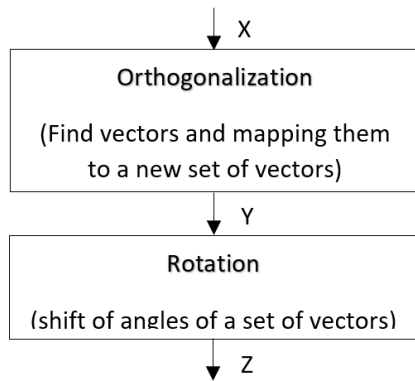


Figure 5. Block diagram of the two-step geometrical approach of the proposed method.

IV. SIMULATION

The effectiveness of the proposed method in the separation of mixed signals is evaluated in this section. Mixed signals are generated by selecting the speaker from the NTT database [22]. The test data of this database is male and female speakers from different countries. Synthetic room impulse responses are made through the Roomsim toolbox [23]. The dimensions of the simulated room are $6.25 \times 4 \times 2.5 \text{ m}^3$ and simulated room was used for the experiments is noise-free. In our experiments, the length of FFT was 256 ms, and the window shift length was 128 ms. The separation quality of the algorithm is measured using the proposed method in [24]. In this method, the separated signals can be decomposed into 3 components:

$$y_p = y_{q_{target}} + y_{q_{interf}} + y_{q_{artif}} \quad (21)$$

where $y_{q_{target}}$ denotes the expected target source after separation with deformation (e.g., filtering or gain), $e_{q_{interf}}$ represents interference caused by unwanted sources, and $e_{q_{artif}}$ refers to artifacts produced by the separation algorithm. The source-to-distortion (SDR), source-to-interference (SIR), and source-to-artifact (dB) ratios are measured as follows:

$$SDR = 10 \log_{10} \frac{\|y_{q_{target}}\|^2}{\|y_{q_{interf}} + y_{q_{artif}}\|^2} \quad (22)$$

$$SIR = 10 \log_{10} \frac{\|y_{q_{target}}\|^2}{\|y_{q_{artif}}\|^2} \quad (23)$$

$$SAR = 10 \log_{10} \frac{\|y_{q_{target}} + y_{q_{interf}}\|^2}{\|y_{q_{artif}}\|^2} \quad (24)$$

According to the literature [26], this toolkit includes SDR, SIR and SAR. The residual crosstalk from other sources, distortion of the target source, spatial distortion, and artifacts can be explained by these criteria; higher numbers are indicative of better performance.

To examine the algorithm performance, in Table 1 the average SDR, SIR and SAR improvement of the proposed algorithm is compared with the Nesta et al [25] and Reju et al. [7] algorithms. The algorithm from Nesta et al. emphasizes its proper performance on short data and the algorithm from Reju et al. performs well and uses its previous data to improve the separation quality but is challenged by short data and a decrease in separation quality. This challenge is seen in short data for all algorithms that use their previous data to improve separation quality.

TABLE I. THE COMPARISON OF THE PROPOSED METHOD 2 SPEAKERS AND 2 MICROPHONES.

	Reju et al.	Nesta et al.	Proposed method
SIR(dB)	6.2	7.4	7.8
SAR(dB)	5.9	7.2	7.4
SDR(dB)	3.5	4.4	4.3

To study the effect of microphone spacing, another experiment was conducted under similar conditions. The difference is that the location of the microphone arrays is considered at five different distances (5, 10, 20, and 40 cm). At all distances, the microphones have a fixed center. If the distance between the microphones is the D variable, the simulated room is considered as Fig. 6.

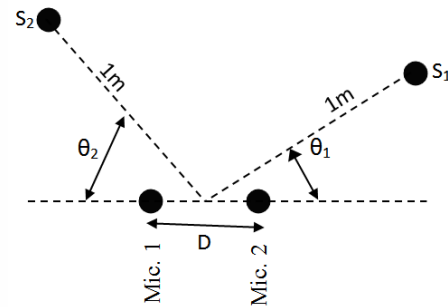


Figure 6. The Synthetic room impulse response.

The results of this experiment are presented in Table 2. It is clear that the efficiency of the separator algorithms and the distance between the two microphones are directly related. The reason is that as the microphones get closer, the impulse response between each source and the various microphone is very similar. As the distance between the microphones increases, the inter-angle of the mixing vectors increases and the efficiency of the separation algorithm also increases. But from a certain distance, this increase in separation algorithm performance will stop. When the value $\theta_1 + \theta_2$ is 180, the microphones are collinear and the efficiency of the separator algorithms is reduced. In this case, the proposed algorithm has the highest increase performance gain compared to other mentioned separation algorithms. When $\theta_1 =$

170. $\theta_2 = 80$ or $\theta_1 = 80$. $\theta_2 = 170$, the proposed algorithm has the worst performance compared to mentioned algorithm. In this experiment, short length data is used and the improvement of the algorithm performance in these conditions is significant.

TABLE II. STUDY THE EFFECT OF MICROPHONE SPACING

Distance between the microphones (cm)		5cm	10cm	20cm	40cm
The SDR average improvement of $\theta_1 = 170$. $\theta_2 = 80$ and $\theta_1 = 80$. $\theta_2 = 170$	Reju et al.	-0.19	-0.02	0.01	0.03
	Nesta et al.	-0.18	-0.03	0.02	0.03
	Proposed method	-0.18	-0.02	0.42	0.51
The SDR average improvement of $\theta_1 = 0$. $\theta_2 = 180$ and $\theta_1 = 0$. $\theta_2 = 180$	Reju et al.	-0.01	0.01	0.03	0.08
	Nesta et al.	0.01	0.02	0.04	0.09
	Proposed method	0.26	0.61	0.93	1.09

To show the increase in performance in short-length data, the efficiency of the algorithm is shown in Fig. 7 by averaging the output value of the algorithm for 10 fixed data while decreasing their length. It can be seen that the algorithm has much better performance in short data. Works [7] and [25] have a high computational time and therefore are not suitable for use in online systems [25]. Moreover, all the tests in Matlab on a 64-bit Intel Core i7-4511U with a 108GHz CPU and their running time of the algorithm was less than 7 seconds and respectively the running time of [25] and [7] were more than 13 and 11 seconds.

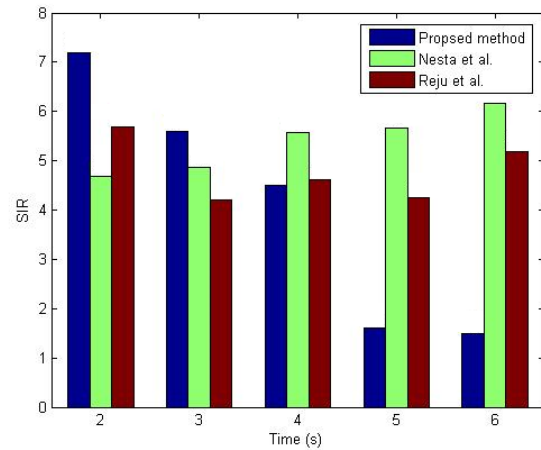
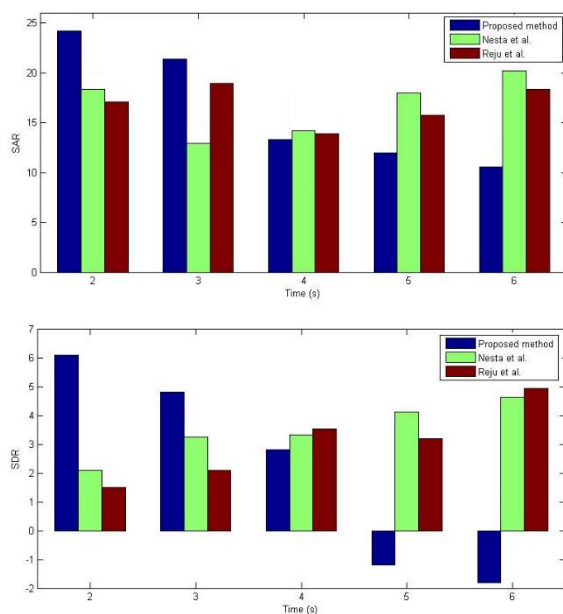


Figure 7. Demonstrate the efficiency of the algorithm in separating short data.

V. CONCLUSION

In this paper, we proposed a new frequency domain BS S method for differentiation of two speech signals captured concurrently by two microphones. This method operates via a two-stage geometrical algorithm applied to discrete frequency bins. In the first stage, observations will be orthogonalized, and in the second stage the orthogonalized data will be rotated. Validity and performance of the proposed method was assessed using three efficiency evaluation criteria and by a number of numerical simulations.

REFERENCES

- [1] Zohrevandi, M., Setayeshi, S., Rabiee, A. et al. Blind separation of underdetermined Convolutional speech mixtures by time-frequency masking with the reduction of musical noise of separated signals. *Multimed Tools Appl* **80**, 12601–12618 (2021). <https://doi.org/10.1007/s11042-020-10398-3>
- [2] D. Nion, K. N. Mokios, N. D. Sidiropoulos, and A. Potamianos, "Batch and adaptive parafac-based blind separation of convolutivespeech mixtures," *IEEE Transactions on Audio, Speech, and Language Processing*, vol. 18, no. 6, pp. 1193–1207, 2010.
- [3] K. Abed-Meraim, Y. Xiang, J. H. Manton, and Y. Hua, "Blind source separation using second-order cyclostationary statistics," *IEEE Transactionson Signal Processing*, vol. 49, no. 4, 2001.
- [4] J. F. Cardoso and B. H. Laheld, "Equivariant adaptive source separation," *IEEE Transactions on Signal Processing*, vol. 44, no. 12, pp. 3017–3030, 1996.
- [5] A. Dapena and L. Castedo, "Stochastic gradient adaptive algorithms for blind source separation," *Signal processing*, vol. 75, no. 1, pp. 11–27, 1999.
- [6] N. Delfosse and P. Loubaton, "Adaptive blind separation of independent sources: a deflation approach," *Signal processing*, vol. 45, no. 1, pp. 59–83, 1995.
- [7] V. G. Reju, S. N. Koh, and Y. Soon, "An algorithm for mixing matrix estimation in instantaneous blind source separation," *Signal cessing*, vol. 89, no. 9, pp. 1762–1773, 2009.
- [8] L. Parra and C. Spence, "Convolutional blind separation of non-stationarysources," *IEEE Transactions on Speech and Audio Processing*, vol. 8, no. 3, pp. 320–327, 2000.

- [9] J. Kocinski, "Speech intelligibility improvement using convolutive blind source separation assisted by denoising algorithms," *Speech Communication*, vol. 50, no. 1, pp. 29–37, 2008.
- [10] J. Liu, J. Xin, and Y. Qi, "A dynamic algorithm for blind separation of convolutive sound mixtures," *Neurocomputing*, vol. 72, no. 1, pp. 521–532, 2008.
- [11] M. Handa, T. Nagai, and A. Kurematsu, "Frequency domain multichannel speech separation and its applications," *IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP)*, vol. 5, pp. 2761–2764, 2001.
- [12] Y. Zhou and B. Xu, "Blind source separation in frequency domain," *Signal processing*, vol. 83, no. 9, pp. 2037–2046, 2003.
- [13] Sawada, H., Mukai, R., Araki, S., & Makino, S., "A robust and precise method for solving the permutation problem of frequency-domain blind source separation," *IEEE Transactions on Speech and Audio Processing*, vol. 12, no. 5, pp. 530-538, 2004.
- [14] L. Drude and R. Haeb-Umbach, "Integration of Neural Networks and Probabilistic Spatial Models for Acoustic Blind Source Separation," in *IEEE Journal of Selected Topics in Signal Processing*, vol. 13, no. 4, pp. 815-826, Aug. 2019, doi: 10.1109/JSTSP.2019.2912565.
- [15] Sahar Zakeri, Masoud Geravanchizadeh, Supervised binaural source separation using auditory attention detection in realistic scenarios, *Applied Acoustics*, Volume 175, 2021, 107826, ISSN 0003-682X, <https://doi.org/10.1016/j.apacoust.2020.107826>.
- [16] Albataineh, Z., Salem, F.M. A RobustICA-based algorithmic system for blind separation of convolutive mixtures. *Int J Speech Technol* (2021). <https://doi.org/10.1007/s10772-021-09833-z>
- [17] Mei, Alfred Mertins, Fuliang Yin, Jiangtao Xi, Joe F. Chicharo, Blind source separation for convolutive mixtures based on the joint diagonalization of power spectral density matrices, *Signal Processing*, Volume 88, Issue 8, 2008, Pages 1990-2007, ISSN 0165-1684, <https://doi.org/10.1016/j.sigpro.2008.02.003>.
- [18] Rabiee, A., Setayeshi, S., Lee, S., CASA: Biologically Inspired Approaches for Auditory Scene Analysis. *Natural Intelligence*, Volume 1, Issue 2, (2012)
- [19] Peng, T., Chen, Y. & Liu, Z. A Time–Frequency Domain Blind Source Separation Method for Underdetermined Instantaneous Mixtures. *Circuits Syst Signal Process* **34**, 3883–3895 (2015). <https://doi.org/10.1007/s00034-015-0035-3>
- [20] D. Wang and J. Chen., 'Supervised speech separation based on deep learning: an overview', arXiv preprint arXiv, pp. 1708-07524 (2017)
- [21] Mukai, Ryo, Shoko Araki, Hiroshi Sawada, and Shoji Makino. "Evaluation of separation and dereverberation performance in frequency domain blind source separation." *Acoustical Science and Technology*, vol. 25, no. 2, pp. 119-126, 2004.
- [22] Cardoso J. F., "Blind signal separation: statistical principles". *Proc. IEEE*, vol. 86, no 10, pp. 2009-2025, 1998.
- [23] NTT Database. Available online: <http://www.ntt-at.com/product> (accessed on 4 December 2017).
- [24] E. Vincent, R. Gribonval and C. Fevotte, "Performance measurement in blind audio source separation," in *IEEE Transactions on Audio, Speech, and Language Processing*, vol. 14, no. 4, pp. 1462-1469, July 2006, doi: 10.1109/TSA.200.5.858005.
- [25] F. Nesta, P. Svaizer and M. Omologo, "Convolutive BSS of Short Mixtures by ICA Recursively Regularized Across Frequencies," in *IEEE Transactions on Audio, Speech, and Language Processing*, vol. 19, no. 3, pp. 624-639, March 2011, doi: 10.1109/TASL.2010.2053027.

Mahbanou Zohrevandi received the B.Sc. degree in Computer Engineering (Hardware) from Shahid Beheshti University in 2004. He received the M.Sc. degree in Computer Engineering (Computer Systems Architecture) from Department of Computer Engineering, Science and Research Branch, Islamic Azad University (SRBIAU), Tehran, Iran in 2008 and 2021. Her research interests are Array Signal Processing and Blind Source Separation.

Saeed Setayeshi

Department of Medical Radiation Engineering, Faculty of Energy Engineering and Physics, Amirkabir University of Technology, Tehran, Iran

Azam Rabiee is assistant professor at Department of Computer Science, Dolatabad Branch, Islamic Azad University, Isfahan, Iran from 2012. In addition, she was Machine Learning researcher at KAIST Institute for Artificial Intelligence (KI4AI), Korea Advanced Institute of Technology (KAIST), South Korea from 2017 to 2019.



Midia Reshadi is currently an Assistant Professor in the Computer Engineering Department at Science and Research Branch of Azad University since 2010. His research interest is Network-on-Chip including Performance and Cost Improvement in Topology, Routing, and Application-Mapping design levels of various types of NoCs such as 3D, Photonic and Wireless. Recently, he has started carrying out research in NoC based Deep Neural Network Accelerators and Silicon Interposer-Based NoC with his team consists of M.Sc. and Ph.D. students.