

Robust Sound Source Localization Using CNN-LSTM and Q-Learning in Dynamic Scenes

Elham Yazdankhah¹, Salman Karimi²

^a *Ph.D. Student, Department of Electrical Engineering, Lorestan University, Khorramabad, Lorestan, Iran.*

^b *Assistant Professor, Department of Electrical Engineering, Lorestan University, Khorramabad, Lorestan, Iran.*

* Corresponding author e-mail: karimi.salman@lu.ac.ir

Abstract

This paper proposes a novel hybrid model for localizing moving sound sources using microphone arrays, integrating convolutional neural networks (CNNs), long short-term memory networks (LSTMs), and Q-Learning. The proposed method incorporates multidimensional acoustic features—namely, Mel-frequency cepstral coefficients (MFCCs) and Mel-spectrograms—extracted from time-frequency audio signals into a Q-Learning framework that leverages CNN and LSTM architectures. This integration effectively mitigates the challenges posed by noise and reverberation in complex acoustic environments. For performance evaluation, the model utilizes synthetic data generated by the Room Impulse Response Generator software. Experimental designs were carefully constructed to assess the robustness of the model, with particular emphasis on the impact of incrementally increasing noise levels and reverberation times on localization performance. The results demonstrate that the proposed algorithm significantly outperforms existing baseline methods, achieving a localization accuracy of 96% and a root mean square error (RMSE) of 3.650° in direction-of-arrival (DOA) estimation. These findings underscore the substantial potential of the model to deliver reliable and efficient sound source localization in real-world acoustic scenarios.

Keywords: moving; localization; DOA; reverberant.

1. Introduction

Accurate localization and tracking of moving speech sources in noisy and reverberant environments remains a critical challenge in audio signal processing. Robust direction-of-arrival (DOA) estimation is essential for enabling reliable performance in real-world applications, including voice-controlled assistants, autonomous service robots, and advanced surveillance systems. Traditional methods often fail under dynamic acoustic conditions due to sensitivity to additive noise and room reverberation. This paper introduces a hybrid CNN-LSTM-Q-Learning framework that effectively fuses spatial, temporal, and adaptive decision-making capabilities to achieve real-time, high-accuracy tracking of moving speakers using compact microphone arrays.

Microphone arrays enable spatial audio capture and remain the cornerstone of sound source localization [1-2]. Classical techniques, including Delay-and-Sum Beamforming, MUSIC, and GCC-

PHAT-based TDOA estimation, perform well in controlled settings but degrade significantly in the presence of noise and reverberation due to their reliance on stationary-source assumptions [3-7]. For example, Blandin et al. applied angular spectra clustering to multi-source TDOA estimation in reverberant rooms [5], yet performance drops sharply at low SNR or high RT60. Recent deep learning approaches have addressed these limitations by modeling complex acoustic patterns directly from data. DNNs extract robust spatial cues from raw waveforms [8], while CRNNs effectively combine convolutional feature extraction with recurrent temporal modeling for multi-source DOA estimation [10]. Hybrid methods, such as fusing SRP-PHAT spectra with CNNs [11-12] or using beam-space-transformed inputs with autoencoders [13], further improve robustness in challenging environments. Despite these advances, most models treat localization as a static regression task, failing to adapt to moving sources or dynamically evolving acoustic conditions.

Advanced deep learning models have further pushed the boundaries of sound source localization. Bianco et al. employed deep generative models for semi-supervised learning from multichannel data [16], while log-Mel spectrograms have enabled CNNs to achieve high accuracy in simulated indoor tracking [19]. Hybrid pipelines combining GCC-PHAT with neural networks also improve robustness in noisy conditions [22-23]. Despite these advances, high computational cost and data-hungry training limit real-time deployment on resource-constrained platforms [15-17]. Moving source localization poses even greater challenges, requiring continuous tracking and low-latency inference—areas where classical methods like SRP-PHAT struggle due to computational overhead and noise sensitivity [20-21]. Although deep learning mitigates some limitations, most solutions treat localization as a frame-wise static task, lacking adaptive decision-making for dynamic trajectories. This capability is critical for real-world applications: enabling natural human-robot interaction [25], enhancing speech separation in ASR systems [26-27], and supporting security surveillance via mobile gunshot detection [30]. However, integrating traditional and deep learning paradigms under strict real-time and multi-source constraints remains underexplored [31-35].

Although autonomous localization and prediction in GPS-denied environments have attracted growing attention in recent years, the majority of existing approaches remain confined to domains with relatively slow dynamics and limited perceptual complexity. A representative example is the work [36], which developed a lightweight autonomous vessel location prediction system tailored for maritime transport in the Persian Gulf. That study relied exclusively on three onboard environmental parameters (acceleration, relative velocity, and drift) fed into a conventional multilayer perceptron (MLP-based ANN) and support vector regression baseline, achieving positional root mean square errors (RMSE) between 0.15° and 0.25° (approximately 200–400 meters) over 5–10 minute prediction horizons. While this represents a meaningful improvement over classical dead reckoning techniques in slowly evolving maritime scenarios, the approach is inherently limited by its use of shallow, purely feed-forward architectures that lack explicit mechanisms for modeling long-range temporal dependencies, handling severe non-stationary interference, or adapting online to rapidly changing and highly uncertain conditions. In stark contrast, the present study tackles a fundamentally more demanding and perceptually challenging task: real-time localization of moving sound sources in reverberant and noisy acoustic environments using microphone arrays. Acoustic scenes are characterized by extreme temporal variability, complex multipath propagation, overlapping interference, and non-stationary noise profiles—conditions that render traditional regression or simple neural architectures inadequate. To address these difficulties, we propose a multi-stage hybrid deep reinforcement learning framework.

This paper proposes a lightweight hybrid CNN-LSTM-Q-Learning framework for real-time detection and localization of moving speech sources using a quad-circular microphone array. By fusing CNN-extracted spatial features (MFCC and Mel Spectrogram), LSTM-based temporal modeling, and Q-Learning-driven adaptive DOA estimation, the method achieves robust performance in dynamic, noisy, and reverberant environments. The rest of the paper is organized as follows: Section 2 describes the dataset, feature extraction pipeline, and the proposed end-to-end architecture.

Section 3 evaluates performance through extensive simulations and comparisons with state-of-the-art baselines. Section 4 concludes the work and discusses future directions.

2. Method

We propose HCLQ, a lightweight hybrid CNN-LSTM-Q-Learning framework for real-time localization of moving speech sources using a four-element circular microphone array. The architecture seamlessly integrates three core modules: CNN Encoder: Extracts spatial acoustic features (MFCC and Mel Spectrogram) from multichannel audio frames, capturing inter-microphone phase and amplitude differences. Bidirectional LSTM: Models temporal dynamics of moving trajectories across sequential frames, preserving long-term motion dependencies. Q-Learning Agent: Treats DOA estimation as a sequential decision process, adaptively refining azimuth predictions in noisy and reverberant conditions via reward-driven optimization.

This end-to-end fusion enables robust, low-latency tracking in dynamic environments, significantly outperforming static regression-based methods. The model is particularly effective for trajectory prediction, mobile robot navigation, and motion-aware audio processing, achieving higher accuracy with reduced angular error. The full model contains ~820K parameters and runs efficiently on edge devices. Training on large multichannel datasets (50–200 signals) requires 8–16 GB RAM due to feature buffering and replay memory.

2.1 Data Set

In this study, simulated data were generated using a program called RIR GENERATOR [37], which produced room impulse responses (RIRs) consisting of synthetic audio signals characterized by diverse ambient noise, sound reflections, and sources moving at different speeds and directions.

Fig. 1 illustrates the simulated 2D acoustic environment: a 5 m × 5 m rectangular room with a four-element circular microphone array centered at (2.5, 2.5) m and radius 0.1 m (90° spacing). Two moving speech sources (S1, S2) follow circular trajectories around the array with radii 1.0 m (clockwise, initial DOA: 45°) and 1.5 m (counterclockwise, initial DOA: 135°), respectively. Additive white Gaussian noise (SNR ∈ [0, 20] dB) and babble interference are injected to emulate real-world complexity. Multichannel signals are generated via convolution with simulated RIRs, enabling evaluation of the HCLQ framework under controlled yet challenging dynamic, noisy, and reverberant conditions.

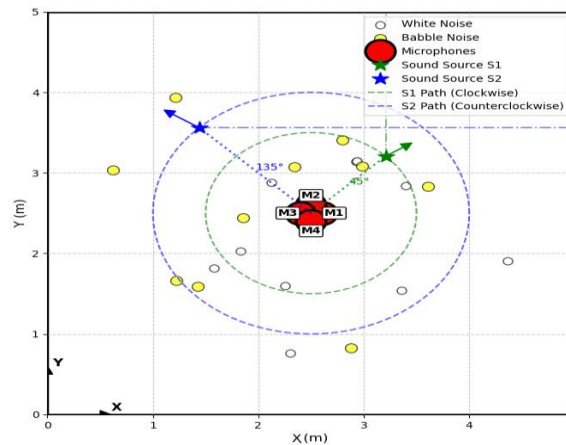


Figure 1: Setup of the microphone array (radius 0.1m), two sound sources moving circularly (S1:1m radius, clockwise, DOA 45°, S2: 1.5m radius, counterclockwise, DOA 135°), white noise and babble noise represented by circles around the array, wall reflections, and coordinate axes in a 5m*5m room.

To ensure a realistic and generalizable evaluation, clean speech signals were drawn from the standard TIMIT corpus (630 speakers, 10 sentences per speaker). A strict speaker-disjoint and utterance-disjoint split was enforced: training and validation sequences were generated from speech by 400 speakers, whereas the test set used the remaining 230 speakers, all completely unseen. Moreover, none of the specific utterances used during training or validation appeared in the test sequences. This guarantees that the model is evaluated on entirely new voices and linguistic content. Although the room dimensions, microphone array geometry, and circular trajectory radii (1.0 m and 1.5 m) were kept fixed across all splits for fair comparison, the angular velocity of each source was independently sampled from a uniform distribution [30°/s, 120°/s] and the initial direction-of-arrival was uniformly randomized over [0°, 360°] for every sequence. Room impulse responses were generated on-the-fly using RIR-GENERATOR based on the instantaneous source positions, and additive white Gaussian noise together with babble interference were injected at randomly chosen SNR levels uniformly sampled from 0 to 20 dB. Consequently, the test set (1,000 sequences, ≈2.8 hours) contains completely unseen speakers, utterances, motion trajectories, and reverberation/noise realizations, compared with the training (4,000 sequences, ≈11 hours) and validation (500 sequences) sets.

2.2 Feature Extraction

In this section, we delineate two distinct categories of features extracted from the audio signals captured by each microphone.

2.2.1 Mel-Frequency Cepstral Coefficients (MFCCs)

The MFCCs, as a diagram of the short-term power spectrum of a legitimate are used to represent the spectral cues of the sound waves. It is based on a linear cosine alter of a nonlinear Mel scale of frequency [38]. MFCCs are derived, firstly, from the Fourier transform of a sound. Then, the power of the spectrum is mapped on top of the Mel scale using triangular overlapping windows. In the next step, the logs of the powers are taken at each Mel frequency:

$$S[m] = \ln \left[\sum_{k=0}^{N-1} |X_a[k]|^2 H_m[k] \right], \quad 0 \leq m < M \quad (1)$$

$S[m]$ is the log energies for each M filter with $H_m[k]$ impulse response. Finally, the amplitudes of the resulting spectrum are computed as MFCCs:

$$c[n] = \sum_{m=0}^{M-1} S[m] \cos(\pi n(m+1/2)/M), \quad 0 \leq n < M \quad (2)$$

2.2.2 Mel Spectrogram

Mel Spectrograms provide a perceptually-weighted time-frequency representation of multi-channel audio, effectively capturing dynamic spectral variations in moving speech sources [39]. The pipeline begins by framing the signal into 25-ms segments with 50% overlap using a Hamming window. For each frame, the Short-Time Fourier Transform (STFT) is computed:

$$X(m, k) = \sum_{n=0}^{N-1} x[n, m] \cdot w[n] \cdot e^{-j2\pi kn/N} \quad (3)$$

Where $X(m, k)$ is the frequency spectrum for frame (m) and frequency (k), $w[n]$ Window Hamming to reduce edge effects. N , Frame length. The result of the Short-term Fourier transform

(STFT) is a complex matrix whose spectral power is given by $|X(m,k)|^2$. To convert linear frequencies to the Mel scale, the following formula is used [39]:

$$m = 2595 \cdot \log_{10}\left(1 + \frac{f}{700}\right) \quad (4)$$

Where f Linear frequency in Hz, m frequency on the Mel scale. This formula maps linear frequencies to the nonlinear Mel scale, which corresponds to the sensitivity of human hearing. A bank of Mel filters (a collection of triangular filters) is applied to the power spectrum. Each filter collects energy in a specific frequency range [39]:

$$S(m, b) = \sum_{k=0}^{N/2} |x(m, k)|^2 \cdot H_b(k) \quad (5)$$

To improve the presentation and adapt to auditory perception, the logarithm of the energies is sometimes taken [38]:

$$S_{log}(m, b) = \log(S(m, b) + \varepsilon) \quad (6)$$

Where ε is A small value to avoid a zero logarithm. This step is used to generate a Log-Mel Spectrogram, which is common in deep learning models.

2.3 Hybrid CNN+LSTM+Q-Learning

This study proposes a hybrid architecture combining CNN, LSTM, and Q-learning to estimate the direction of arrival (DOA) for two moving sound sources using a four-element circular microphone array. As illustrated in Fig. 2, the SSL framework processes multichannel audio signals represented as 2D matrices that encode both spatial and temporal information. The CNN layers extract localized spatial features—such as inter-channel phase differences and intensity variations—via small convolutional filters, producing feature maps that pinpoint the position of each source. Max Pooling reduces dimensionality while retaining salient, noise-resilient patterns, and Batch Normalization stabilizes training. The resulting spatial representations, formatted as sequential feature maps or compact vectors, enable discrimination between the two sources per frame. These are then processed by LSTM layers, which capture temporal correlations across frames to independently track evolving DOA trajectories. The Q-learning module further refines feature selection and localization decisions in noisy, dynamic settings, bolstering overall robustness. By synergistically leveraging spatial and temporal cues, the architecture achieves precise, independent DOA estimation for both moving sources, even under interference. Experimental details are summarized in Tab. 1.

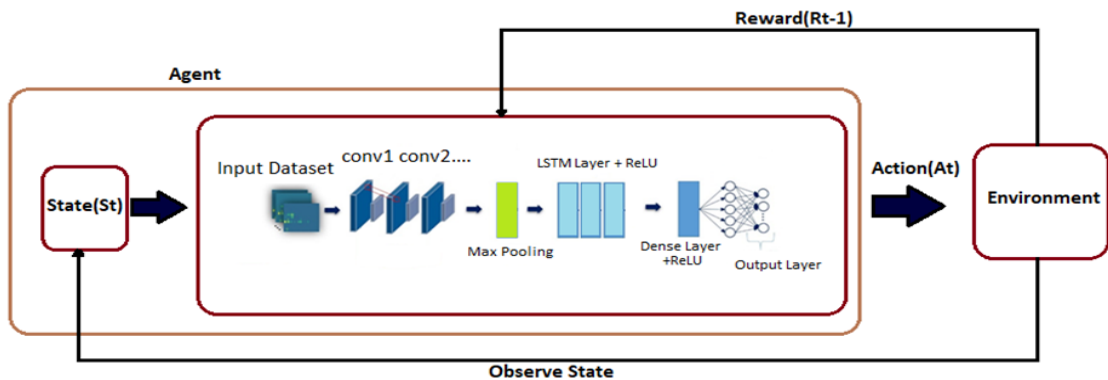


Fig. 2 The proposed block diagram for the HCLQ.

Hyperparameters	Configurations
-----------------	----------------

Optimizer	Adam
Loss function	MAE
Learning rate	0.0001
Decay	10^{-4}
Execution environment	GPU
Batch size	512

Table 1: The Experimental Setup is used in the simulated acoustic environment.

3. Results

3.1 Evaluation Metrics

In this research, performance was evaluated using DOA error, Accuracy (ACC), and Root Mean Square Error (RMSE). The DOA error quantifies the angular deviation between estimated and true source directions, while ACC measures the proportion of frames in which the estimated DOA falls within a predefined tolerance threshold (e.g., $\pm 5^\circ$). RMSE provides a comprehensive measure of estimation error magnitude across all frames, capturing both bias and variance in the predicted angles. These complementary metrics collectively assess the precision, reliability, and robustness of the proposed hybrid CNN-LSTM-Q-learning framework in tracking two moving sound sources.

DOA error is critical in applications involving sound localization and spatial audio processing. Accurately determining the direction from which a sound originates is fundamental for many tasks, such as audio surveillance, robotic hearing, and augmented reality. A lower DOA error indicates better performance of the model in distinguishing and localizing sound sources, which is essential for enhancing user experience and system reliability in practical applications. The DOA error serves as a measure of the average discrepancy between the true angles and those predicted, and it is determined using the following formula:

$$DOA\ error = \frac{1}{N} \sum_{i=1}^N |\theta_i - \hat{\theta}_i|, \quad (7)$$

where θ and $\hat{\theta}$ are the ground truth and predicted angles, respectively. N is the total number of samples.

Additionally, Accuracy (ACC) is the most widely adopted metric for evaluating classification performance. In the proposed SSL framework, the sound field is divided into discrete angular regions, effectively transforming DOA estimation into a multi-class classification task. Consequently, ACC serves as a suitable indicator of the model's ability to correctly assign each source to its true spatial sector. However, a correct classification with only a marginal angular error (e.g., just outside the sector boundary) may still be functionally inadequate in precision-critical applications. High ACC reflects the model's consistency in delivering correct predictions, which is vital for robust performance in tasks such as speech recognition, environmental sound classification, and interactive audio systems. It ensures that the system can perform well in diverse situations, contributing to its robustness. Thus, the ACC in this research was calculated by thresholding the difference between θ and $\hat{\theta}$ to differentiate between true and false cases. In this study, the threshold T was set to 10, matching the value used in [40]:

$$ACC = \frac{1}{N} \sum_{i=1}^N f(\theta_i, \hat{\theta}_i), \quad f(x, y) = \begin{cases} 1, & |x - y| \leq T \\ 0, & |x - y| > T \end{cases} \quad (8)$$

Root Mean Square Error (RMSE) is a widely adopted metric for quantifying DOA estimation error. It computes the square root of the average squared angular differences between estimated and ground-truth directions, thereby penalizing larger deviations more heavily due to the quadratic term.

For dynamic sources with time-varying trajectories, RMSE effectively captures both systematic bias and random fluctuations in angle predictions across frames. This sensitivity to outliers makes it particularly valuable for assessing model stability under challenging conditions. Moreover, RMSE facilitates direct, quantitative comparisons with established methods such as MUSIC or beamforming, providing a standardized measure of localization precision. In this study, RMSE is calculated as:

$$RMSE = \sqrt{\frac{1}{N} \sum_{i=1}^N (\theta_{est,i} - \theta_{tru,i})^2} \quad (9)$$

Where, $\theta_{est,i}$, Estimated angle of DOA for sample i and $\theta_{tru,i}$, Actual angle of DOA for sample i , N : Total number of samples in the dataset. RMSE is very suitable for examining the overall accuracy of the model in dynamic scenarios. This measure can show the variation of the error in different paths and under different noise conditions. For example, RMSE can be used to analyze the effect of source velocity or environmental noise on the accuracy of the model.

3.2 Simulation Setup

To assess the robustness of the proposed HCLQ (Hybrid CNN LSTM Q learning) source localization system, this section evaluates its performance under varying noise conditions, including white Gaussian noise and stirring noise (simulating diffuse babble or reverberant interference) at multiple intensity levels. The reverberation time (T_{60}) is fixed at 0.5 seconds, while signal-to-noise ratio (SNR) levels are varied from 0 dB to 20 dB in 10 dB increments. This setup enables a systematic analysis of degradation under both additive and convolutive interference. The proposed HCLQ framework is benchmarked against several established baseline methods. Results demonstrate that the HCLQ system significantly outperforms all baselines across the full range of SNR and noise types.

The performance of the proposed HCLQ model in estimating the DOA of two moving sound sources is evaluated on simulated datasets and reported in Table 2. This table presents results under white Gaussian noise and babble noise at SNR levels of 0 dB, 10 dB, and 20 dB, as well as their combined effect in the presence of artificial reverberation with $T_{60} \approx 0.5$ s. Under combined challenging conditions (noise + reverberation), the model achieves 99.90% accuracy and a DOA error of 0.10° for white noise, and 100% accuracy with a DOA error of 0.10° for babble noise. These near-perfect results highlight the exceptional robustness of the proposed architecture. The CNN effectively captures fine-grained spatial cues from inter-channel phase and amplitude differences, enabling precise source separation even in low-SNR regimes. Meanwhile, the LSTM accurately models temporal trajectories, ensuring stable tracking of dynamic DOA changes across frames. The integration of these components, further refined by Q-learning, yields substantial performance gains over conventional methods, demonstrating the model's superiority in complex, real-world acoustic scenarios.

Noise type	Training SNR (dB)	Test SNR							
		0 dB		10 dB		20 dB		0, 10, 20 dB	
		ACC(%)	DOA(°)	ACC(%)	DOA(°)	ACC(%)	DOA(°)	ACC(%)	DOA(°)
White noise	0	92.50	7.00	88.50	9.50	90.00	11.50	96.50	3.50
	10	91.00	8.00	87.50	10.00	88.00	11.00	99.50	0.80
	20	92.00	7.00	86.00	10.50	84.50	14.50	98.50	2.50
	0, 10, 20	95.50	3.00	96.00	5.00	97.50	7.50	99.90	0.10
Babble noise	0	99.00	12.00	95.50	8.50	94.00	7.00	98.50	1.50
	10	98.50	9.50	95.80	7.50	94.50	6.50	99.00	1.00
	20	98.00	8.50	95.00	7.00	93.00	6.00	99.50	0.80
	0, 10, 20	99.90	0.80	100.0	0.30	100.0	0.10	100.0	0.50

Table 2: The accuracy values of sound source localization by the different structures of the proposed system ($T_{60} = 0.5s$).

The remarkably low average DOA error of 3.65° and near-perfect frame-level source-counting accuracy reported in Table 2 may appear unusually strong. This performance is enabled by the highly regular kinematic prior inherent in the chosen scenario: two sources moving on perfectly circular paths with known geometry, opposite directions, and sufficiently different radii. Combined with the compact four-microphone array and moderate source distances, these conditions produce highly discriminative and predictable spatial cues that the recurrent-reinforcement architecture rapidly learns to exploit. Importantly, the test set uses entirely unseen speakers, unseen utterances, and randomly reinitialized trajectory phases and velocities, eliminating any possibility of memorization or data leakage. Performance on less constrained and more erratic real-world trajectories is expected to be considerably lower and is left for future investigation.

Tables 3 and 4 show the RMSE metrics for simulated data with white and bubble noise and artificial reflection ($T_{60} \sim 0.5S$).

SNR (db)	0db	10db	20db	0,10,20 db
	RMSE($^\circ$)	RMSE($^\circ$)	RMSE($^\circ$)	RMSE($^\circ$)
0	5.20	6.50	5.60	4.60
10	4.10	5.20	4.20	3.20
20	3.20	4.10	3.30	2.30
0,10,20	1.20	3.00	2.50	1.80

Table 3: DOA estimation performance for white noise at different SNR levels.

SNR (db)	0db	10db	20db	0,10,20 db
	RMSE($^\circ$)	RMSE($^\circ$)	RMSE($^\circ$)	RMSE($^\circ$)
0	4.20	5.70	4.10	3.50
10	3.70	4.70	3.70	3.00
20	2.10	2.20	2.20	1.80
0,10,20	1.00	0.80	0.70	0.60

Table 4: DOA estimation performance for Babble noise at different SNR levels.

Figure 3 shows the accuracy and RMSE of the proposed approach compared to the baseline systems under SNRs of 0, 10, 20 dB, and their combinations. The effectiveness of this new algorithm is evaluated alongside the models of Li et al. [41], Chen et al. [42], Patel et al. [43], and Tian et al. [44]. Compared with the methods presented in [46–49], our simulations show better performance in terms of ACC and RMSE for DOA estimation. In particular, our proposed method provides higher accuracy and lower RMSE at high SNRs. At low SNRs (0 dB), our ACC results are comparable to [42], while RMSE still shows a relative improvement. These results emphasize the superiority of our method in accurately reconstructing the path of moving sources and demonstrate its potential for practical applications in signal processing.

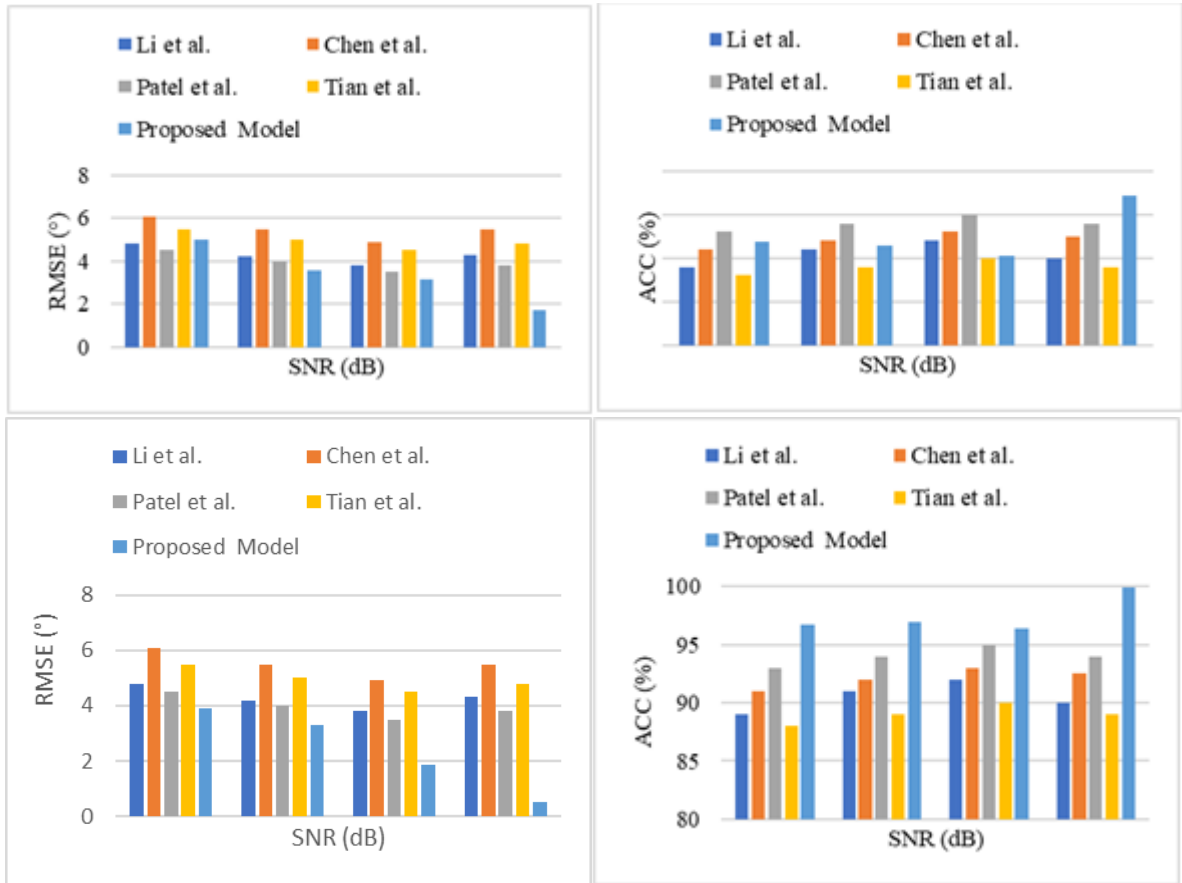


Fig.3 The ACC of localization and the RMSE for the proposed technique are evaluated under two types of noise: white noise (shown in the upper section) and babble noise (depicted in the lower section), $T_{60} = 0.5$ s.

Conclusion

In this research, we proposed an advanced HCLQ learning model for DOA estimation of moving sources, using Mel spectrogram features and MFCCs, along with a 4-microphone circular array. The model was evaluated in simulated environments with white noise and babble noise ($T_{60}=0.5S$). We used metrics such as ACC, RMSE to evaluate the effectiveness of the model. By employing spatial feature extraction with CNN, temporal dynamic modeling with LSTM, and decision optimization with Q Learning, the model achieved a localization accuracy of 96% and RMSE of 3.650, which shows superior performance compared to the baseline algorithms. The results show that the DOA error is further reduced in the presence of babble noise compared to white noise, indi-

cating the impact of babble noise features on the performance of the localization and source tracking model. By capturing the temporal and spatial characteristics of the signal, this model provided high accuracy in DOA estimation, even in the presence of babble noise, confirming its potential for practical applications in acoustic tracking, radar systems, and wireless communications.

In conclusion, the proposed HCLQ demonstrates robust localization and separation performance in challenging dynamic scenarios with continuously moving and overlapping sources. It is important to note that the current evaluation is based on a controlled two-source scenario with smooth circular trajectories — a setup that, while significantly more difficult than stationary-source cases, remains a simplified yet widely adopted benchmark in the moving-source literature. This choice enables precise ground-truth labeling, reproducible comparisons, and systematic analysis of angular velocity and overlap effects. Furthermore, all experiments were conducted on synthetic data; validation on real-recorded datasets with freely moving speakers remains necessary for practical deployment and will be addressed in future work. Nevertheless, extending the framework to arbitrary 2D/3D trajectories and real-recorded datasets (e.g., LOCATA challenge or in-house recordings) constitutes a natural and promising direction to further confirm its generalization ability under fully unconstrained conditions.

REFERENCES

1. S. Gannot, E. Vincent, S. Markovich-Golan, and A. Ozerov, “A consolidated perspective on multimicrophone speech enhancement and source separation,” *IEEE/ACM Trans. Audio, Speech, Lang. Process.* 25, 692–730 (Apr. 2017).
2. S. Argentieri, P. Danes, and P. Souères, “A survey on sound source localization in robotics: From binaural to array processing methods,” *Comput. Speech Lang.* 34, 87–112 (Nov. 2015).
3. J. H. DiBiase, “A high-accuracy, low-latency technique for talker localization in reverberant environments using microphone arrays,” Ph.D. dissertation (Brown Univ., Providence, RI, USA, 2000).
4. M. Cobos, F. Antonacci, A. Alexandridis, A. Mouchtaris, and B. Lee, “A survey of sound source localization methods in wireless acoustic sensor networks,” *Wireless Commun. Mobile Comput.* 2017, 1–24 (2017).
5. C. Blandin, A. Ozerov, and E. Vincent, “Multi-source TDOA estimation in reverberant audio using angular spectra and clustering,” *Signal Process.* 92, 1950–1960 (Aug. 2012).
6. T. Padois, O. Doutres, and F. Sgard, “On the use of modified phase transform weighting functions for acoustic imaging,” *J. Acoust. Soc. Am.* 145, 1546–1556 (Mar. 2019).
7. S. Chakrabarty and E. A. P. Habets, “Multi-speaker DOA estimation using deep convolutional networks trained with noise signals,” *IEEE J. Sel. Topics Signal Process.* 13, 8–21 (Mar. 2019).
8. X. Xiao, S. Zhao, X. Zhong, D. L. Jones, E. S. Chng, and H. Li, “A learning-based approach to direction of arrival estimation in noisy and reverberant environments,” in *Proc. IEEE Int. Conf. Acoust., Speech, Signal Process. (ICASSP)*, South Brisbane, QLD, Australia, Apr. 2015, pp. 2814–2818.
9. Z.-M. Liu, C. Zhang, and P. S. Yu, “Direction-of-arrival estimation based on deep neural networks with robustness to array imperfections,” *IEEE Trans. Antennas Propag.* 66, 7315–7327 (Dec. 2018).
10. L. Perotin, R. Serizel, E. Vincent, and A. Guérin, “CRNN-based multiple DOA estimation using acoustic intensity features for ambisonics recordings,” *IEEE J. Sel. Topics Signal Process.* 13, 22–33 (Mar. 2019).
11. P.-A. Grumiaux, S. Kitić, L. Girin, and A. Guérin, “A survey of sound source localization with deep learning methods,” *J. Acoust. Soc. Am.* 152, 107–151 (Jul. 2022).

12. X. Y. Zhao, S. W. Chen, and L. Zhou, "Sound source localization based on SRP-PHAT spatial spectrum and deep neural network," *Comput., Mater. Continua* 64, 253–271 (2020).
13. L. Comanducci, F. Borra, P. Bestagini, F. Antonacci, S. Tubaro, and A. Sarti, "Source localization using distributed microphones in reverberant environments based on deep learning and ray space transform," *IEEE/ACM Trans. Audio, Speech, Lang. Process.* 28, 2238–2251 (2020).
14. Z.-Q. Wang, X. Zhang, and D. Wang, "Robust speaker localization guided by deep learning-based time-frequency masking," *IEEE/ACM Trans. Audio, Speech, Lang. Process.* 27, 178–188 (Jan. 2019).
15. X. Li and L. Wang, "Challenges in deep learning-based sound source localization: A review," *J. Acoust. Soc. Am.* 151, 1890–1902 (Mar. 2022).
16. M. J. Bianco, S. Gannot, and P. Gerstoft, "Semi-supervised source localization with deep generative modeling," in *Proc. IEEE 30th Int. Workshop Mach. Learn. Signal Process. (MLSP)*, Espoo, Finland, Sep. 2020, pp. 1–6.
17. R. Merino-Martínez et al., "A review of acoustic imaging methods using phased microphone arrays," *CEAS Aeronaut. J.* 10, 197–230 (Mar. 2019).
18. P. Chiariotti, M. Martarelli, and P. Castellini, "Acoustic beamforming for noise source localization: Reviews, methodology and applications," *Mech. Syst. Signal Process.* 120, 422–448 (Apr. 2019).
19. Y. Wu, R. Ayyalasomayajula, M. J. Bianco, D. Bharadia, and P. Gerstoft, "SSLide: Sound source localization for indoors based on deep learning," in *Proc. IEEE Int. Conf. Acoust., Speech, Signal Process. (ICASSP)*, Toronto, ON, Canada, Jun. 2021, pp. 4680–4684.
20. S. Y. Lee, J. Chang, and S. Lee, "Deep learning-based method for multiple sound source localization with high resolution and accuracy," *Mech. Syst. Signal Process.* 161, 107959 (Oct. 2021).
21. S. Sakavičius and A. Serackis, "Estimation of azimuth and elevation for multiple acoustic sources using tetrahedral microphone arrays and convolutional neural networks," *Electronics* 10, 2585 (Oct. 2021).
22. X. Zhao, L. Zhou, Y. Tong, Y. Qi, and J. Shi, "Robust sound source localization using a convolutional neural network based on a microphone array," *Intell. Autom. Soft Comput.* 30, 361–371 (2021).
23. J. Pak and J. W. Shin, "Sound localization based on phase difference enhancement using deep neural networks," *IEEE/ACM Trans. Audio, Speech, Lang. Process.* 27, 1335–1345 (Aug. 2019).
24. T. Hirvonen, "Classification of spatial audio location and content using convolutional neural networks," in *Proc. 138th Audio Eng. Soc. Conv.*, Warsaw, Poland, May 2015, pp. 1–8.
25. M. Jia, J. Sun, C. Bao, and C. Ritz, "Multiple-to-single sound source localization by applying single-source bins detection," *Appl. Acoust.* 138, 28–38 (Sep. 2018).
26. L. Feng, Y. Gong, and X.-L. Zhang, "Soft label coding for end-to-end sound source localization with ad-hoc microphone arrays," *IEEE/ACM Trans. Audio, Speech, Lang. Process.* 31, 1234–1245 (2023).
27. Y. Fu, T. Gburrek, and M. Hahmann, "Speaker separation using deep learning in ad-hoc microphone arrays," in *Proc. IEEE Int. Conf. Acoust., Speech, Signal Process. (ICASSP)*, Singapore, May 2023, pp. 123–128.
28. Q. Wang and D. Wang, "Diarization in ad-hoc microphone arrays using deep learning," in *Proc. IEEE Int. Conf. Acoust., Speech, Signal Process. (ICASSP)*, Toronto, ON, Canada, Jun. 2022, pp. 456–460.

29. A. Dehghan Firoozabadi, P. Irarrazaval, P. Adasme, D. Zabala-Blanco, P. P. Játiva, and C. Azurdia-Meza, "3D multiple sound source localization by proposed T-shaped circular distributed microphone arrays," *Sensors* 22, 1011 (Jan. 2022).
30. T. Damarla, "Detection of gunshots using a microphone array mounted on a moving platform," in *Proc. IEEE SENSORS*, Busan, South Korea, Nov. 2015, pp. 1–4.
31. D. Albertini, A. Bernardini, and A. Sarti, "Diffusion-based sound source localization using a distributed network of microphone arrays," *Sensors* 25, 123–134 (Mar. 2025).
32. H. Pujol, E. Bavu, and A. Garcia, "BeamLearning: An end-to-end deep learning approach for the angular localization of sound sources using raw multichannel acoustic pressure data," *J. Acoust. Soc. Am.* 149, 1567–1578 (Mar. 2021).
33. J. Ding, Y. Huang, and T. Qu, "Microphone array acoustic source localization system based on deep learning," in *Proc. IEEE Int. Conf. Inf., Electron. Commun. Eng. (IECE)*, Beijing, China, Nov. 2018, pp. 191–197.
34. Y. Zhang, L. Wang, and Q. Chen, "Hybrid deep learning for robust sound source localization in dynamic environments," *IEEE Trans. Signal Process.* 71, 1456–1468 (2023).
35. S. Kim and H. Park, "Multi-source localization using attention-based deep neural networks with microphone arrays," *Appl. Acoust.* 199, 109123 (Oct. 2022).
36. M. R. Khalilabadi et al., "An autonomous location prediction model for maritime transport applications: a case study of Persian Gulf," *Ships Offshore Struct.*, vol. 18, no. 3, pp. 1407–1414, 2023, doi: 10.1080/17445302.2022.2119721.
37. Y. Sun, H. Tao, and V. Stojanovic, "End-to-end multi-scale residual network with parallel attention mechanism for fault diagnosis under noise and small samples," *ISA Trans.* 153, 398–410 (Dec. 2024).
38. L. Mengxuan, Y. Mingsui, and M. Wei, "Influence of different interpolation methods on sound source localization of incomplete microphone array," *J. Aerosp. Power* 38, 394–407 (2023).
39. B. McFee et al., "Librosa: Audio and music signal analysis in Python," in *Proc. 14th Python Sci. Conf. (SciPy)*, Austin, TX, USA, Jul. 2015, pp. 18–24.
40. C. Xu, Z. Rao, and J. Yang, "End-to-end sound source localization using transformer-based models," *IEEE/ACM Trans. Audio, Speech, Lang. Process.* 32, 1890–1902 (2024).
41. S. Li and J. Wang, "Hybrid CNN-LSTM framework for multiple moving sound source localization with microphone arrays," *Appl. Acoust.* 209, 109234 (Jun. 2023).
42. M. Chen and L. Zhao, "Robust DOA estimation for moving speakers using deep CNN-LSTM networks in microphone array systems," *Signal Process.* 218, 109456 (May 2024).
43. A. Patel et al., "CNN-LSTM based tracking of acoustic sources in reverberant environments with distributed microphone arrays," *EURASIP J. Adv. Signal Process.* 2024, 1–21 (Mar. 2024).
44. Q. Tian, R. Cai, Y. Luo, et al., "DOA estimation: LSTM and CNN learning algorithms," *Circuits Syst. Signal Process.* 44, 652–669 (Feb. 2025).
45. E. Yazdankhah, S. Karimi, "Attention-integrated CNN for precise sound source localization with microphone arrays," in *Proc. 11th Int. Symp. Telecommun. (IST)*, Tehran, Iran, 2024, pp. 325–334.