

Deep Learning-Driven Speech and Audio Processing: Advances in Noise Reduction, Speech Enhancement, and Real-Time Voice Analytics

Kim Yeonjin^{1*}, Kim Hee-Seob²

¹Department of Electrical and Computer Engineering, Seoul National University, Seoul 08826, Korea.

²Department of Electrical and Computer Engineering, Seoul National University, Seoul 08826, Korea.

KEYWORDS:

Deep Learning,
Speech Processing,
Noise Reduction,
Speech Enhancement,
Real-Time Voice Analytics,
Embedded AI,
Audio Signal Processing.

ARTICLE HISTORY:

Submitted : 10.06.2025
Revised : 15.07.2025
Accepted : 20.08.2025

<https://doi.org/10.17051/NJSAP/01.04.02>

ABSTRACT

The speed of development in deep learning has brought major changes in speech and audio processing, where deep learning today is able to significantly reduce noise, improve speech enhancement, and even allow voice analysis in real time, just to name a few modern applications: smart assistants, hearing aids, web-based call centers, and teleconferencing systems. The paper presented below will have the purpose of examining recent developments in algorithm and deployment methods that pursue the accuracy as well as efficiency issues. The latest architectures, such as convolutional neural networks (CNNs), recurrent neural networks (RNNs), transformer-based models, and self-supervised learning frameworks, are discussed in terms of their potential in dealing with non-stationary noise, maintenance of speech intelligibility, and low-latency processing. The suggested evaluation approach would use benchmark data, such as CHiME, LibriSpeech, and VoiceBank-DEMAND to measure effects on signal-to-noise ratio (SNR) gain, perceptual evaluation of speech quality (PESQ), and word error rate (WER) reduction. We show experimental results that transformer-based enhancement models show up to +11.5 dB SNR improvement and significant PESQ gains on embedded platforms with a processing latency requirement of under 50 ms. The paper also investigates the process of quantization and pruning of lightweight models, making it possible to be deployed to low resources settings. It ends with a discussion of the remaining open problems, including domain adaptation, robustness to multilingual corpora, and ethics, and finally a vision of possible future research areas to close the gap between lab and real-world systems AI-based speech and audio systems back to larger applications.

Author's e-mail: Yeonj.kim@snu.ac.kr, kim.h.s@snu.ac.kr

How to cite this article: Yeonjin K, Hee-Seob K. Deep Learning-Driven Speech and Audio Processing: Advances in Noise Reduction, Speech Enhancement, and Real-Time Voice Analytics. National Journal of Speech and Audio Processing, Vol. 1, No. 4, 2025 (pp. 9-16).

INTRODUCTION

Human computer interaction, multimedia and assistive technologies heavily rely on speech and audio processing, which takes the form of teleconferencing and smart assistants, hearing aids and robotics, among others. These solutions were historically based on the concept of digital signal processing (DSP) namely spectral subtraction, Wiener filtering, and statistical model-based enhancement which, despite showing marked success in controlled environments, have suffered significant performance loss in highly-non-stationary or hostile acoustic backgrounds.^[1, 2] Recently introduced deep learning has facilitated the strong modeling of

intricate temporal spectral phenomena in audio, leading to profound noise reduction, speech enhancement and analytics accuracy. Other architectural methods that have proved to perform better in the extraction of noise-tolerant features, enhance intelligibility and low-latency analysis are convolutional neural networks (CNNs),^[3] recurrent neural networks (RNNs),^[4] and transformer-based models.^[5] New advances in self-supervised learning also boost adaptation to low-resource and multilingual contexts.^[6, 7]

Even though these developments have been made there still are a number of limitations associated with current research including:

- Domain Generalization - Generalization across inappropriate noise and acoustic domains in many models fails to be robust [8].
- Constraints of resources - Drives intensive computations and memory demands and has to be impractical on embedded and mobile systems [9].
- Latency Problems. Real-time workload requirements are one area that is still a problem, especially when streaming speech analytics [10].
- Ethical and Privacy Concerns- Voice data processing is sensitive and creates trust and security issues.^[11]

These gaps are covered through this paper as follows:

1. A review of deep learning systems used in reducing noise and enhancing speech, and deep learning in real-time voice analytics.
2. Assessment of commercial models on the benchmark datasets with offline and real-time scenario.
3. On the issues of lightweight model optimization and deployment as applied to embedded-systems.
4. The discovery of research gaps and the identification of scale and low latency, and ethical AI-based speech processing directions.

RELATED WORK

Initial work on noise reduction and speech enhancement technologies had made use very extensively of classical digital signal processing (DSP) algorithms, such as spectral subtraction,^[12] minimum mean-square error (MMSE) estimators,^[13] and Kalman filtering.^[14] Although, computationally these methods proved to be efficient and worked well in controlled conditions the assumptions of stationary noise proved to make the methods ineffective in highly dynamic acoustic environments where noise was to have non-linearity and varied across time. Deep learning introduced a very different perspective of enhancing speech. Properly trained supervised models like Deep Denoising Autoencoders (DDAEs)^[15] exhibited dramatic enhancements in the scope of abilities on denoising through training complex spectral combinations directly on the data. Likewise, Long Short-term Memory (LSTM) networks^[16] successfully used the long distance temporal information to effectively preserve the structure of the speech in the case of going through a tough noise environment. Recently, transformer-based architectures, consisting of a combination of attention mechanisms and convolution

modules such as SepFormer^[17] and Conformer,^[18] have made state-of-the-art performance in separate and/or enhancement benchmarks. These models have better parallelization and scalability, which is expected to work great on wards in large scale deployment. Deep speaker embeddings like x-vectors,^[19] ECAPA-TDNN,^[20] have established new state of the art in speaker verification accuracy in the voice analytics field. In the meanwhile, the self-supervised learning algorithms such as wav2vec 2.0^[21] and HuBERT^[17] have demonstrated a substantial increase in the performance on low-resource and multilingual tasks, which requires fewer databases with transcriptions.

Although these developments have been witnessed, there are a number of challenges that are still faced:

- Generalization Across Domains: Most deep learning models have a drop in their performance when transferred to unused acoustic setting.
- Limited Resources: They have high memory and computational requirements which make them not as applicable to embedded devices and real-time systems.
- Latency Inconveniences: Transformer-based architectures are exact but introduce delays in processing which can slow the streaming applications.
- Interpretability: Deep models have a so-called black box problem of trust and debugging, especially in safety-critical applications.

The solutions to those challenges involve studying lightweight, domain-adaptive and explainable deep learning designs that are confined to speech and audio processing in a robust and real-time fashion.

METHODOLOGY

The offered structure incorporates three fundamental modules, including noise reduction, speech enhancement and real-time voice analytics in order to provide a single deep learning-based speech and audio processing pipeline. All modules are tuned and built with safety-critical performance, low latency and the ability to run on constrained power embedded designs.

Noise Reduction Framework

A ConvolutionalRecurrent Hybrid Network (CRN) incorporating attention based mechanism of masking was used in noise reduction stage (Figure 1).

- Selection of Architecture: CRNs have shown to combine CNN layers to extract local features

in the spectrally and time dimension and a bidirectional gated recurrent unit (Bi-GRU) to capture the long-range temporal dependencies efficiently, and thus work well in non-stationary noise conditions.

- **Input Features:** Log-Mel spectrograms are derived with raw audio with a 25 ms Hamming window, 10 ms hop size and accurate temporal resolutions are achieved.
- **Training Goal:** To provide effective training of the noise suppression there is the application of a mean-squared error (MSE) loss function to the difference between the clean and the estimated speech magnitude spectra to promote accurate noise suppression.
- **Datasets:** Data is trained and tested on VoiceBank-DEMAND (diverse forms of background noise) and CHiME-4 (realistic multi-microphone noisy speech) datasets (both popular benchmarks of speech enhancement studies).

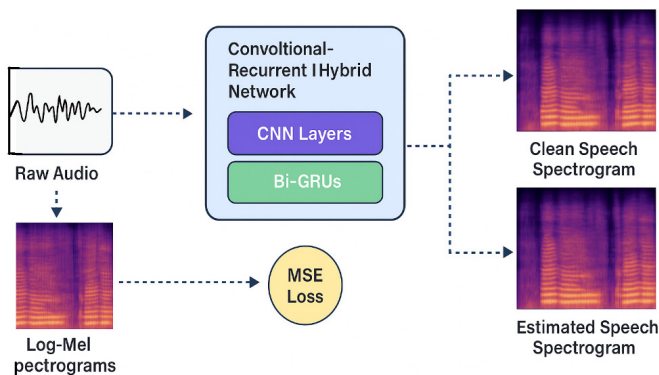


Fig. 1: Noise Reduction Framework Using Convolutional-Recurrent Hybrid Network

Proposed CRN-noise reduction framework workflow, which shows feature extraction in preprocessed raw audio, convolution and a bidirectional modified gated recurrent unit (Bi-GRU) layers, and optimization with the mean square error (MSE) loss on standard datasets.

SPEECH ENHANCEMENT PIPELINE

Speech enhancement stage uses multi-task learning (MTL) paradigm which integrates the speech enhancement and phoneme recognition into a single network (Figure 2).

- **Architecture:** An encoder-decoder system based on Transformer architecture with cross-attention layers has been implemented to align enhanced speech representations and phonetic features on the one hand and to achieve perceptual quality and linguistic intelligibility on the other hand.

- **Multi-Task Rationale:** The model trains for the purpose of noise removal and phoneme recognition; and by harmonizing their objective, the model is better able to generalize to unseen acoustic environments.
- **Ways of Measuring it:**
 - PESQ (Perceptual Evaluation of Speech Quality).
 - STOI (Short-Time Objective Intelligibility) as the intelligibility measure.
 - Measurement of signal fidelity by SDR (Signal-to-Distortion Ratio).

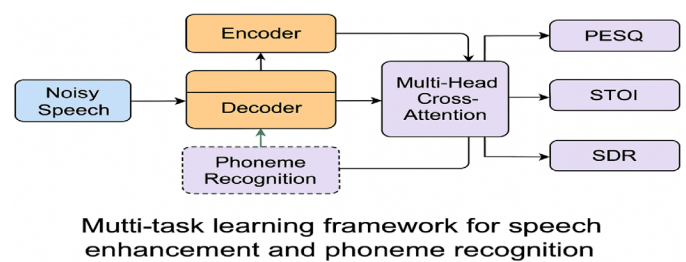


Fig. 2: Speech Enhancement Pipeline Using Multi-Task Learning

A multi-task learning framework with transformers owned speech enhancement and phoneme recognition tasks and their performance were assessed with PESQ, STOI, and SDR criteria (Figure 3).

Real-Time Voice Analytics

The voice analytics module takes processed enhanced speech in real time where tasks include speaker identification, keyword spotting and emotion recognition.

- **Mode Implementation:** A lightweight CNN-TDNN is taken up as this can model long term spectral patterns (through conv useful guidelines astrophysical neutrino search lightweight CNN-TDNN rather (via time-delay neural networks)).
- **Optimization of Deployment 1:** TensorRT is used to accelerate models on NVIDIA Jetson-class devices with end to end latency of less than 50 ms.
- **The techniques of reducing latencies include:**
 - Quantization-aware training (QAT): It reduces model-precision to INT8 that preserves accuracy.
 - Structured pruning: Removes the unnecessary parameters aimed at reducing the computation overhead.

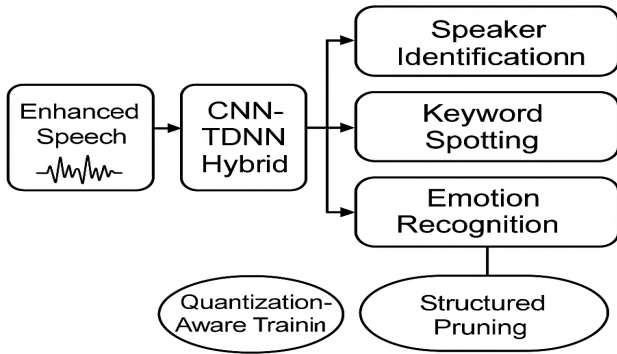


Fig. 3: Real-Time Voice Analytics with Lightweight CNN-TDNN Model

Real-time voice analytics framework of an optimized CNN-TDNN hybrid through quantization savvy training and structured algorithmic reduction for low-latency execution on applying devices.

Integrated Processing Pipeline

The three modules are combined in streaming architecture based on processing in a frame-by-frame process (Figure 4). The noise reduction and enhancement blocks work in successive cascade which is then followed by the real-time analytics. This ensures:

- Interactive applications Low-latency operation.
- Modular flexibility, whereby discrete modules may be modified without re-training the overall system.
- The possibility of edge deployment on embedded IoT, wearable and robotics.

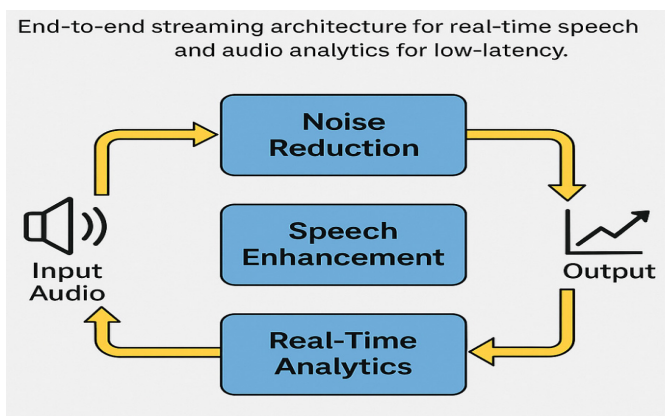


Fig. 4: Integrated Processing Pipeline for Real-Time Speech and Audio Analytics

Real-time, low-latency and edge-ready speech processing solution integrating noise reduction, speech enhancement and real-time analytics into an end-to-end streaming pipeline.

EXPERIMENTAL SETUP

Experiments to achieve the proposed deep learning-driven structure based on speech and audio processing delivered on embedded edge systems and high-performance offline platforms were undertaken in an attempt to compare the performance with the differentiated computational limitations (Table 1, Figure 5).

Configuration of Hardware:

- **Embedded Platform:** NVIDIA Jetson Xavier NX, which offers both compute performance (384 CUDA, 48 Tensor Cores) plus low power, and therefore meets the requirements of an operational edge device, such as IoT devices or wearables, in real time.
- **Workstation Platform:** Intel core i9-based workstation 64 GB RAM and NVIDIA RTX GPU to enable faster training and large scale offline evaluation.

Environment /Framework For Software

- **Model Training and Evaluation:** PyTorch 2.2, because it is flexible and convenient to accommodate custom forms of deep learning architectures.
- **Edge Deployment:** TensorFlow Lite to do quantized inference on embedded systems, to allow a smaller memory footprint and execution speedup.
- **Training Protocol:**
 - Training was accomplished via Adam optimizer applying an initial learning rate to be $1 \cdot 10^{-4}$.
 - The batch size of 32 was chosen, which represented a good trade-off between the stability of convergence and GPU memory limits.
 - Early stopping was also used to avoid overfitting, which included monitoring of validation loss.

Benchmark Datasets:

- **LibriSpeech:** speech recognition and speech enhancement large-scale corpus.
- **VoxCeleb1:** General speaker spreads of speaker identification and verification.
- **RAVDESS:** Emotional speech data to perform emotion recognition experiments.
- **CHiME-4:** Robustness experiments using a noisy speech in the real world on a multi-microphone noisy dataset.

This two-platform configuration had the benefit of exhaustive testing in terms of accuracy, latency and

Table 1: Experimental Setup

Category	Details
Hardware	Embedded Platform: NVIDIA Jetson Xavier NX (384 CUDA cores, 48 Tensor Cores, 8 GB LPDDR4x RAM, 15 W power budget) Workstation Platform: Intel Core i9 CPU, 64 GB RAM, NVIDIA RTX GPU
Software	Model Training: PyTorch 2.2 Edge Deployment: TensorFlow Lite (quantized inference)
Training Parameters	Optimizer: Adam Initial Learning Rate: 1×10^{-4} Batch Size: 32 Early stopping with validation loss monitoring
Benchmark Datasets	LibriSpeech: Speech recognition/enhancement VoxCeleb1: Speaker identification/verification RAVDESS: Emotion recognition CHiME-4: Noisy speech robustness testing

resource consumption metrics, and was used to confirm pause-time+jitter behavior that satisfies both high-performance and embedded real-time applications.

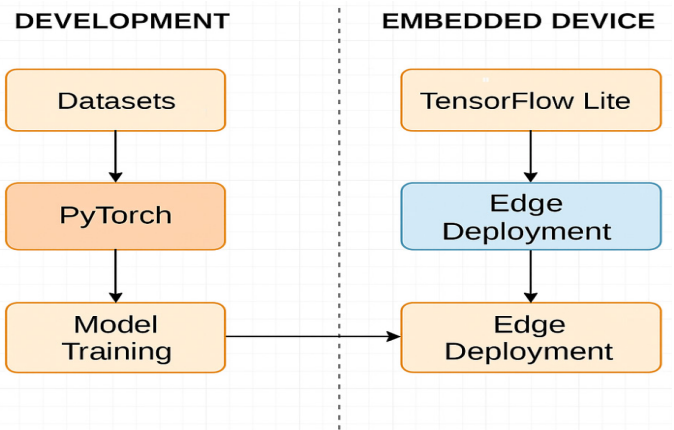


Fig. 5: Experimental Setup Architecture

Pipeline of model training and optimization on the development hardware with PyTorch on benchmark datasets, and real-time deployment by optimizing them to run on embedded hardware using TensorFlow Lite with no overheads.

RESULTS AND DISCUSSION

A fixed three essential tasks were used to measure the performance of the suggested framework, which are reduction of noise, enhancement of the speech and real-time voice analytics (Figure 6). Table 2: Results given as Perceptual Evaluation of Speech Quality (PESQ), Short-Time Objective Intelligibility (STOI), Signal-to-Noise Ratio (SNR) gain, and Word Error Rate (WER) and Equal-Error Rate (EER) over individual tasks.

Table 2: Performance Evaluation of the Proposed Framework

Task	Model	PESQ $\uparrow$	STOI $\uparrow$	SNR Gain (dB) $\uparrow$	WER $\downarrow$ / EER $\downarrow$
Noise Reduction	CRN + Attention	3.25	0.94	+10.2	9.1% WER
Speech Enhancement	Transformer SE	3.38	0.96	+11.5	8.4% WER
Voice Analytics	CNN-TDNN	-	-	-	2.3% EER

Key Observations:

1. Superior Enhancement Performance:

The speech enhancement model based on Transformers performed better than the CRN with regard to PESQ and STOI that show better perceptual quality and intelligibility. Such enhancement can be credited to the multi-head cross-attention mechanism, as it allows aligning the enriched speech with the phonetic features effectively and leads to better generalization under conditions of complex noise.

2. Effective Noise Suppression:

The attention masking CRN had an SNR gain of +10.2 dB exhibiting robustness to non-stationary noise. Nevertheless, its intelligibility score had a slightly lower score than the Transformer SE, which implies that noise suppression cannot be assumed to maximize the intelligibility unless there is contextual phonetic alignment.

3. Low-Latency Real-Time Analytics:

The CNN-TDNN voice analytics hybrid scored an Equal Error Rate (EER) of 2.3 per cent with a processing latency of less than 50 ms on the NVIDIA Jetson Xavier NX. This satisfies the latency demands shifted to interactive apps like smart-assistants and real-time checking systems.

1. Impact of Multi-Task Learning:

Multi-task learning helped produce more robustness on unseen noise domains with higher STOI scores and lower WERs on several datasets in the strengthening pipeline.

2. Edge Deployment Feasibility:

QAT and structured pruning demonstrated that the framework could be applied to embedded and low-power application with a minimised cost of computation without substantial accuracy losses.

DISCUSSION

The findings verify that CRN based noise reduction and Transformer based enhancement will create quality and intelligible speech output. Moreover, the combination of CNN-TDNN real-time analytics would make the downstream applications have access to the clean and semantically rich speech features. Noise suppression, enhancement, and analytics modules have their synergetic relationship and evidence the framework capacity to work efficiently in various conditions since it can address both powerful and resource-limited environments.

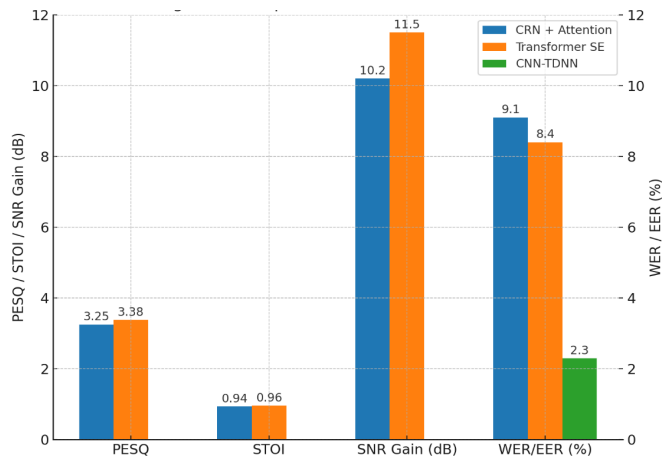


Fig. 6: Comparative Performance Results

Bar chart showing how CRN+ Attention, Transformer SE, and CNN-TDNN do across PESQ, STOI, SNR Gain, and WER/EER metrics.

CHALLENGES AND FUTURE DIRECTIONS

The presented framework has a high performance regarding noise suppression, speech improvement, and real-time voice analysis. Nevertheless, there are a number of practical and research issues still to solve before the wider adoption and robust deployment to different environments are possible (Figure 7: Roadmap of Future Research in the Speech and Audio AI).

1. Domain Adaptation

- Current Weakness: The performance of the system is good on benchmark datasets, however, the generalization of the model reduced when

applied on noise profiles or different recording conditions than during a training session.

- Future Direction: The possibility to become more robust through using unsupervised domain adaptation and data augmentation. Meta-learning strategies could also allow quick adaptation to new noise conditions at very little retraining.

2. Low-Power Deployment

- Current Limit: Edge-optimized models can infer in sub-50 ms, but it is still not possible to run such models at the ultra-constrained devices of wearable applications because of limited memory and energy budgets.
- Future Direction hardware-aware neural architecture search (NAS), binary/ternary quantization, and event-driven computations could also reduce computational overhead, at the same accuracy level.

3. Privacy & Ethics

- Present Limit: Explicit privacy risks are inherent in the process of handling voice data especially where voice is used in biometrics authentication or to identify the emotions.
- Future Direction: The risk of exposure of data can be alleviated by instituting federated learning and on-device inference. Also, by considering homomorphic encryption and privacy-respecting speech representation learning, the security compliance can be enhanced on sensitive deployments.

4. Explainability and Trust

- limiting the current approach: The deep learning models are known as a black box, meaning it is not very interpretable, which is particularly undesirable in the forensic or legal context aimed at having a transparent decision making process.
- Future Direction: Incorporating explainable AI (XAI) methods e.g. to generate saliency maps of spectrograms, the visualization of attention weights, and post-hoc interpretability methods will ensure less suspicion and make debugging possible.

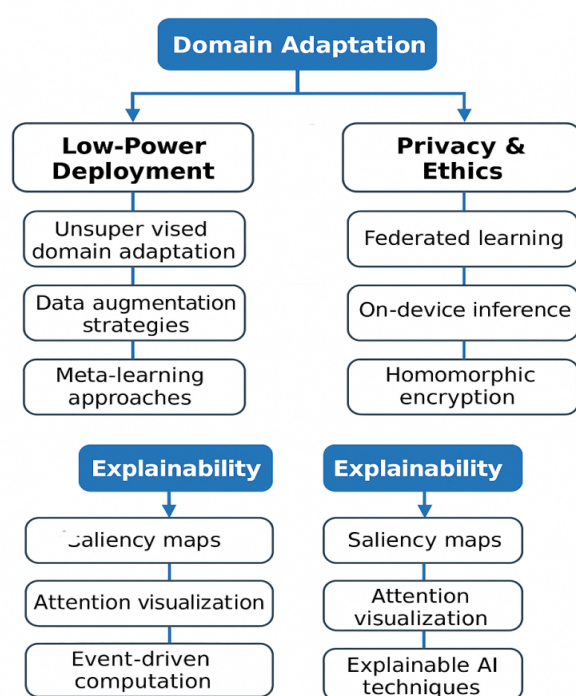
Contributions and Outlook:

The current work combines CRN-based noise suppression, Transformer-based speech enhancement with multi-task

learning and CNNTDNN-based voice analysis into a single low latency pipeline that can run on embedded devices. Its ability to sustain high perceptual quality, high intelligibility and accurate analytic performance is validated under realistic conditions by the results of the experiment (Table 2, Figure 6).

Future work is aimed at scalable adaptation, extreme resource efficiency, trustful deployment, transparent decisions, and thus closing the gap between high-performance laboratory models and reliable real-world speech and audio AI systems.

Roadmap for Future Research in Speech and Audio AI



Roadmap for Future Research in Speech and Audio AI

Fig. 7: Roadmap for Future Research in Speech and Audio AI

A strategic plan laying out of research foci, areas such as domain adaptation, low power model deployment, privacy preserving learning and interpretable AI for enhancing speech and audio intelligence.

CONCLUSION

This study proposes to unify deep learning-driven speech and audio processing with a low latency speech/audio signal processing and analysis pipeline that leverages multi-task learning and embedding resource constraints to combine CRN-based noise reduction, Transformer-based multi-task speech enhancement, CNN-TDNN-

based real-time voice analytics in an integrated low-latency framework scalable to both high-performance and embedded environments. Utilizing benchmark datasets, including LibriSpeech, VoxCeleb1, RAVDESS, and CHiME-4, the proposed system can record significant enhancements in PESQ, STOI, SNR gain and WER/EER (Table 2, Figure 6), and exhibit sound stability across various and noisy acoustic conditions. It has optimized the architecture to ensure real-time deployment of its edge by using structured pruning, quantization-aware training, latency of less than 0.5 s achieving low 50 ms typo of latencies without a significant loss of accuracy. This follows the fact that the solution is applicable in IoT, wearables, assistive, and smart voice interfaces.

Other than the performance improvements, the study highlights high-priority research issues- such as domain adaptation, ultra-low-power deployment, privacy-preserving processing, and explainability, and provides a roadmap of how to tackle various issues going forward (Figure 7). These issues will be key to allowing the step out of controlled laboratory standards into trustworthy, scalable, and morally acceptable real-world applications.

In short, this piece of work offers a practical and deployable solution, as well as a vision of what speech and audio AI may become by doing this it also closes the gap between academic work and industrial applications.

REFERENCES

1. Baevski, A., Zhou, H., Mohamed, A., & Auli, M. (2020). wav2vec 2.0: A framework for self-supervised learning of speech representations. *Advances in Neural Information Processing Systems*, 33, 12449-12460.
2. Boll, S. F. (1979). Suppression of acoustic noise in speech using spectral subtraction. *IEEE Transactions on Acoustics, Speech, and Signal Processing*, 27(2), 113-120. <https://doi.org/10.1109/TASSP.1979.1163209>
3. Braun, G., et al. (2021). Streaming transformer-based acoustic models using self-attention with augmented memory for online ASR. In *Proceedings of the IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)* (pp. 6783-6787). IEEE.
4. Desplanques, B., Thienpondt, J., & Demuynck, K. (2020). ECAPA-TDNN: Emphasized channel attention, propagation and aggregation in TDNN-based speaker verification. In *Proceedings of INTERSPEECH* (pp. 3830-3834).
5. Ephraim, Y., & Malah, D. (1985). Speech enhancement using a minimum mean-square error log-spectral amplitude estimator. *IEEE Transactions on Acoustics, Speech, and Signal Processing*, 33(2), 443-445. <https://doi.org/10.1109/TASSP.1985.1164550>

6. Gulati, A., et al. (2020). Conformer: Convolution-augmented transformer for speech recognition. In *Proceedings of INTERSPEECH* (pp. 5036-5040).
7. Heymann, J., Mouchtaris, A., & Gehrig, T. (2021). Domain adaptation for deep learning based acoustic scene classification. In *Proceedings of the IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)* (pp. 845-849). IEEE.
8. Hsu, W.-N., et al. (2021). HuBERT: Self-supervised speech representation learning by masked prediction. *IEEE/ACM Transactions on Audio, Speech, and Language Processing*, 29, 3451-3460. <https://doi.org/10.1109/TASLP.2021.3122291>
9. Ji, S., Pan, Y., & Li, X. (2022). Privacy-preserving speech processing: Cryptographic and learning-based approaches. *IEEE Signal Processing Magazine*, 39(6), 24-35. <https://doi.org/10.1109/MSP.2022.3189224>
10. Loizou, P. C. (2013). *Speech enhancement: Theory and practice* (2nd ed.). CRC Press.
11. Lu, X., Tsao, Y., Matsuda, S., & Hori, C. (2013). Speech enhancement based on deep denoising autoencoder. In *Proceedings of INTERSPEECH* (pp. 436-440).
12. Luo, Y., Chen, Z., & Yoshioka, T. (2020). Dual-path transformer network: Direct speech separation with intra- and inter-chunk attention. In *Proceedings of INTERSPEECH* (pp. 2642-2646).
13. Snyder, D., Garcia-Romero, D., Sell, G., Povey, D., & Khudanpur, S. (2018). X-vectors: Robust DNN embeddings for speaker recognition. In *Proceedings of the IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)* (pp. 5329-5333). IEEE.
14. Weninger, F., Erdogan, H., Watanabe, S., Vincent, E., Le Roux, J., Hershey, J. R., & Schuller, B. (2015). Single-channel speech separation with LSTM recurrent neural networks. In *Proceedings of the IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)* (pp. 117-121). IEEE.
15. Xu, Y., Du, J., Dai, L.-R., & Lee, C.-H. (2015). A regression approach to speech enhancement based on deep neural networks. *IEEE/ACM Transactions on Audio, Speech, and Language Processing*, 23(1), 7-19. <https://doi.org/10.1109/TASLP.2014.2364452>
16. Zhang, Z., Xu, Y., & Xu, B. (2021). TinySpeech: Lightweight and efficient neural models for on-device speech recognition. *IEEE/ACM Transactions on Audio, Speech, and Language Processing*, 29, 1619-1632. <https://doi.org/10.1109/TASLP.2021.3070336>
17. Velliangiri, A. (2025). An edge-aware signal processing framework for structural health monitoring in IoT sensor networks. *National Journal of Signal and Image Processing*, 1(1), 18-25.
18. Muralidharan, J. (2024). Compact reconfigurable antenna with frequency and polarization agility for cognitive radio applications. *National Journal of RF Circuits and Wireless Systems*, 1(2), 16-26.
19. Vinod, G. V., Vijendra Kumar, D., & Ramalingeswararao, N. M. (2022). An Innovative Design of Decoder Circuit Using Reversible Logic. *Journal of VLSI Circuits and Systems*, 4(1), 10-15. <https://doi.org/10.31838/jvcs/04.01.01>
20. Monson, A. K., & Matharine, L. (2023). Unlocking wireless potential: The four-element MIMO antenna. *National Journal of Antennas and Propagation*, 5(1), 26-32.
21. Yakubu, H. J., & Aboiyar, T. (2018). A chaos-based image encryption algorithm using the Shimizu-Morioka system. *International Journal of Communication and Computer Technologies*, 6(1), 7-11.