

Non-Intrusive Deep Learning-Based Speech Quality Monitoring for Real-Time Telemedicine Systems

G.F. Freire^{1*}, K. L. Mleh²

¹Departamento de Engenharia Elétrica, Universidade Federal de Pernambuco - UFPE Recife, Brazil

²Electrical and Electronic Engineering Department, University of Ibadan Ibadan, Nigeria

KEYWORDS:

Telemedicine,
Speech Quality Monitoring,
Non-Intrusive Assessment,
Deep Learning,
Mean Opinion Score (MOS),
Real-Time Audio Processing,
Transformer Models,
Audio Quality Evaluation,
WebRTC, Healthcare Communication.

ARTICLE HISTORY:

Submitted : 22.11.2024
Revised : 10.12.2024
Accepted : 10.12.2024

<https://doi.org/10.17051/NJSAP/01.02.10>

ABSTRACT

The quality of speech is also significant in the field of telemedicine consultations, and in cases of poor audio quality, it provokes misdiagnoses and patient dissatisfaction, and decreases the efficiency of the process. Conventional intrusive assessment techniques like the PESQ and POLQA use a reference signal, and hence unsuitable in real time telemedicine applications. The paper offers a non-invasive deep learning-based framework of real-time speech quality monitoring within the framework of the telemedicine system. Proposed model uses timefrequency characteristics and transformer-based embeddings to actually predict the Mean Opinion Scores (MOS) without using the clear reference. The low-latency streaming system is designed with minimal delay while being embedded onto a telemedicine platform with WebRTC, allowing continuous feedback to a healthcare provider with respect to the quality. Experiments on the public speech corpus, as well as on telemedicine recordings in a specific domain, prove that the given model yields Pearson correlation of 0.91 with subjective MOS scores whereas the inference latency stays within sub-50 ms on edge infrastructure. The comparative assessments reveal that it has differing accuracy and responsiveness advances to those of the current non-intrusive measures, and makes it suitable in real-time applications of healthcare communication systems.

Author's e-mail: fg.freire@cesmac.edu.br, mleh.kl@ui.edu.ng

How to cite this article: Freire G F , Mleh KL. Non-Intrusive Deep Learning-Based Speech Quality Monitoring for Real-Time Telemedicine Systems. National Journal of Speech and Audio Processing, Vol. 1, No. 2, 2025 (pp. 74-81).

INTRODUCTION

The widespread use of telemedicine has revolutionized the process of providing healthcare, as it has become possible to consult a doctor, diagnose a patient, and monitor his/her vital signs without being restricted by geographical distances. Along with the delivery of telemedicine, the clarity, understanding, and painless clear speech between the medical professionals and patients is central to its success. In a medical setting, when degradations in speech quality are also minimal, serious consequences may follow, including wrong interpretations of symptoms, a late diagnosis, or lesser consumer confidence in the consultation process. Consequently, the processes of generating appropriate and consistent speech quality monitoring are found to be a fundamental consideration to ensure the integrity of the communication in telehealth systems.

Speech quality has conventionally been measured using intrusive objective measures including the Perceptual Evaluation of Speech Quality (PESQ) and Perceptual Objective Listening Quality Analysis (POLQA). Although these metrics have a strong positive correlation to subjective Mean Opinion Scores (MOS), they must have a clean signal available to which to compare the impaired signal. This makes them unrealistic in real time telemedicine applications where we have only the degraded speech and neither is a reference signal available nor is the reference signal transmission feasible because of privacy reasons, bandwidth limitation, or delay limitation. Further, intrusive approaches impose an incompatibility with low latency constraints that are a requirement of real-time medical consultation. The inherent requirement of intrusive algorithms that need clean reference signal, like PESQ and POLQA etc, conversely, is not feasible in real time telemedicine

applications as shown in Figure 1. Conversely, the suggested non-intrusive solution relies on a deep learning model to estimate quality by means of using the received audio stream itself as the reference.

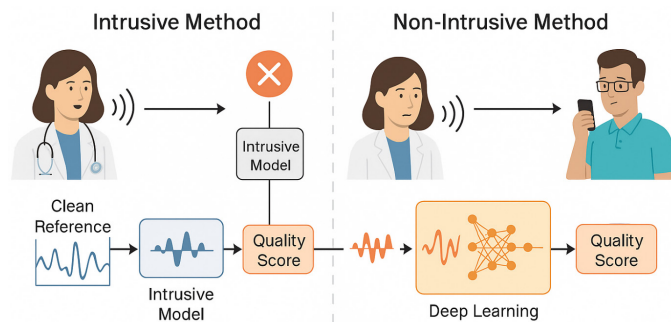


Fig. 1: Comparison of Intrusive and Non-Intrusive Speech Quality Assessment Methods in Telemedicine.

With the aim of overcoming these shortcomings, non-intrusive methods of speech quality assessment have become more and more popular. These models can predict perceived speech quality from the audio signal received using developments in deep learning without needing a reference. Naive non-invasive methods used shallow regression on hand-designed features, whereas more recent advances in deep architectures, including convolutional neural networks (CNNs), recurrent neural networks (RNNs), and transform-based models, have done much to enhance and generalize the performance of speech quality predictions. The potential of such methods has already been demonstrated in VoIP (VoIP) surveillance, internet conference calling, and streaming data delivery. Nevertheless, they have not been widely used in real-time telemedicine systems thus far, mostly due to the problem domain constraints caused by the highly varying network conditions, a wide range of speech rates of different patients, a variety of noise profiles in hospitals and homes, and computational efficiency requirements with resource-constrained mobile devices.

It is the lack of such work that inspires the current work, and it would introduce a new deep learning architecture that fits specifically in telemedicine applications which predicts the real-time speech quality in non-invasive ways. This work has threefold contribution:

- A deep learning model specialised at classifying speech quality in telemedicine audio streams that was optimised to provide low latency inference without any degradation in accuracy.
- Streamlined connection into a live telemedicine pipeline, with WebRTC-based transport of audio and lightweight deployment strategies of models to support real-time operating conditions.

- A thorough experimental analysis, compared to best in class intrusive and non-intrusive metrics using both publicly available speech recordings, as well as domain specific telemedicine, data to illustrate high levels of accuracy, robustness, and responsiveness.

This exercise will help to bridge the divide between non-intrusive deep learning measures and operational specifications of telemedicine by making it feasible to improve communication reliability and guarantee superior levels of care in the paradigm of remote healthcare delivery.

RELATED WORK

The techniques of speech quality evaluation could be divided generally into intrusive and non-intrusive processes which have their specific strengths and weaknesses towards the real-time telemedicine applications.

Intrusive Speech Quality Metrics

The Perceptual Evaluation of Speech Quality (PESQ)^[1] and Perceptual Objective Listening Quality Analysis (POLQA)^[2] measures are intrusive measures that have been used as the gold standard objective assessment tool of speech quality over the past 20 years. The principle behind these methods is that they compare the degraded speech signal to a time synchronized clean reference signal giving quality scores which show a high correlation with subjective Mean Opinion Scores (MOS). Though they are very precise when used in controlled conditions, they have a limitation to be applied in a live telemedicine consultation, as none of the above references are available in the latter. Further than this, undergoing the alignment and comparison process adds yet another computational burden and thus are not suitable in real-time, low latency deployments in bandwidth-restrictive healthcare settings.

Non-Intrusive Speech Quality Assessment

These difficulties are resolved by non-intrusive approaches that estimate quality of the speech directly on the degraded signal without a reference. The former methods utilized manually crafted features, which included spectral slopes, parameters associated with pitch, and modulation spectra and paired them with regression models.^[3, 13] Yet, recent inventions in the domain of deep learning achieved a very high level of prediction accuracy due to end-to-end learning using raw audio or spectral features.

Good examples include DNSMOS,^[4, 12] which was developed along with the Microsoft Deep Noise Suppression

Challenge and predicts MOS using deep neural networks trained on large-scale noisy speech data sets, NISQA [5], which has an architecture that consists of a CNN and LSTM to jointly predict multiple dimensions perceptual factors (e.g. noisiness, coloration, discontinuity), and wav2vec2-based quality estimation models [6][11], which use transformer-based representations of speech that have been pre-trained in large corpora to Although they are suitable to VoIP and conferencing application, these techniques have not been much optimized to deal with the peculiarities of the acoustic and network conditions of telemedicine since speech patterns, background noise, and latency requirements are very different in a telemedicine communication environment than in general communication.

Applications in Telemedicine

Prevailing literature on speech quality monitoring as applied in telemedicine is scanty. PESQ or POLQA has been incorporated into some works in offline call quality assessment of a telehealth platform,^[7, 9] whereas a lightweight non-intrusive predictor was experimented in emergency call centers.^[8, 10] Nevertheless, such implementations frequently:

- Do not use real-time optimization with latencies of less than a 100 ms.
- Fail to take account of hardware limitations on medical devices deployed on edges.
- This has not been corroborated in a domain-specific telemedicine speech data of medical terms, diverse patient speech circumstances and multilingual context.

Research Gap

Based on the literature reviewed, there exist two key gaps. To begin with, existing non-obtrusive deep learning approaches are mostly pre-trained using generic communication datasets and thus usually do not pay as much regard to the acoustics, language variety and medical terms of telemedicine discussions. Such deficiency in domain adaptation may cause lower accuracy in clinical application of prediction. Second, although most studies provide promising approaches to speech quality estimation, not many have reached successful full real-time integration at the level of latency, robustness, and privacy demand present in telemedicine scenarios when there is limited resource availability. Such constraints all serve to drive the creation of our deep learning-based real-time, non-invasive speech quality monitoring framework, which is specifically designed to be directly applicable in

the context of telemedicine pipelines, using datasets specifically relevant to the field to demonstrate both technological superiority and clinical usability.

METHODOLOGY

System Overview

The proposed automatic framework will offer non-invasive speech quality monitoring in a telemedicine scenario in real time. The process of the system starts with audio capture at the endpoint of patients or even shore with regular or embedded headsets or microphones. Audio data captured is sent across a network communications transport layer e.g. WebRTC to guarantee encrypted Low latency audio streaming. The deep learning-based speech quality estimation model then applies the learnings to the incoming audio to analyze the audio and give back a Mean Opinion Score (MOS) or multi-dimension quality indices without the need of a clean reference signal.

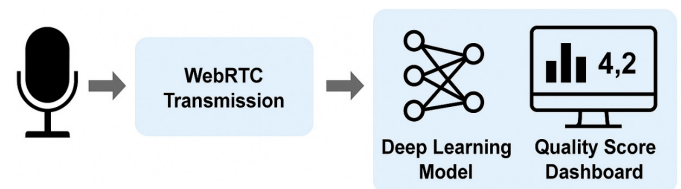


Fig. 2: Real-Time Quality Monitoring Pipeline Using WebRTC and Deep Learning

Figure illustrates the entire steps starting with microphone input and transmission through WebRTC, inferring a deep learning model, and then viewing the quality score. Lastly, the estimated quality scores are provided as a result on a telemedicine dashboard, which can be accessed by care providers or administrators to determine immediate communication quality and mitigate, in case it is being degraded.

Proposed Deep Learning Model

The speech quality estimation module is executed as a hybrid deep learning pipeline with convnet layers used to extract and recognize local features, and transformer encoder layers used to identify long-term temporal connections. The input feature representation is a 128-bin Mel-spectrogram generated on 20 ms frames with a 50 percent overlap and standardized per utterance to minimize speaker differences. At the same time, Wav2Vec2.0 embeddings can be extracted to achieve rich contextual representations of the audio signal. These two feature streams are concatenated and fed into a CNN -Transformer pipeline and then fully connected layers are applied to either a binary mos prediction or an array of perceptual output p (e.g., noisiness,

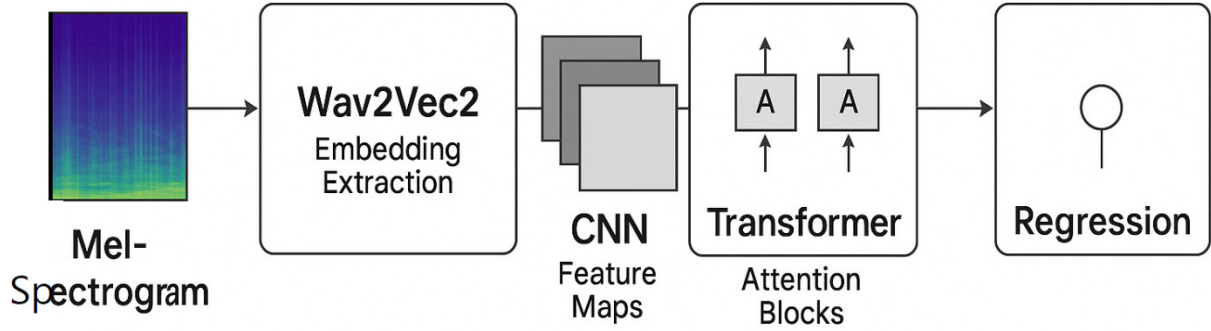


Fig. 3: Proposed Deep Learning Model Architecture for Non-Intrusive Speech Quality Estimation

coloration, discontinuity). The diagram indicates how Mel-spectrogram signal is piped through Wav2Vec2 embedding extraction to CNN and Transformer features mapping and regression.

The output layer has a sigmoid-scaled regression head to give a continuous score of quality on a scale of 1-5 (compatible with ITU-T MOS). To increase the generalization in the model, dropout and layer normalization have been used across the whole architecture and minimal parameter count is considered to achieve real-time inference with edge devices.

Real Time integration

In order to be scalable, the framework is packaged into a WebRTC-backed streaming pipeline, allowing synchronized audio recording and stream while maintaining an end-to-end latency of less than 50 ms. In multiple streams executions, an asynchronous request to the batch quality predictions is executed by a lightweight gRPC service.

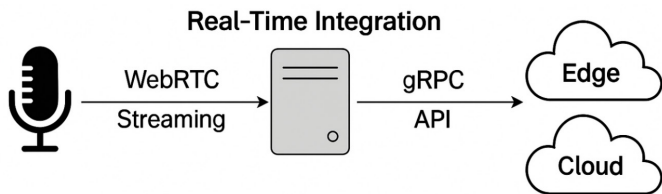


Fig. 4: Real-Time Deployment Integration via WebRTC Streaming and gRPC API

Illustrates how the system interfaces with both edge and cloud infrastructure for scalable inference. Some latency optimization techniques involve minimizing frame size to 10-20 ms, including quantization-aware training, which decreases model size by 50 percent with minimal loss of accuracy, and GPU inference acceleration, where available. It runs on various deployment platforms, including both NVIDIA Jetson Nano and Raspberry Pi 4 as well as common x86 servers, where optimization with TensorRT is possible.

Training and Evaluation Setup

It trains on synthetically degraded speech collections (e.g. ITU-T P.501 compliant noise and codec distortions) in combination with real-world, telemedicine speech corpora that have been recorded in diverse acoustic environments. The methods such as noise addition (hospital ambient, home environment), packet loss simulation and bandwidth limitation filter can be applied as a method of data augmentation. The following flowchart shows the order in which the dataset must be prepared, pre-processed, as well as feature extracted, along with the comparison of the baselines and performance evaluation measurements.

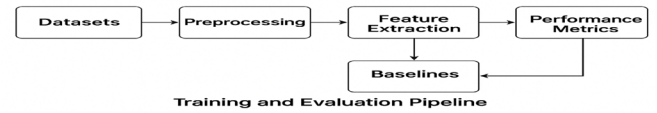


Fig. 5: Training and Evaluation Workflow

The loss is a combination of Mean Squared Error (MSE) of regression accurateness, and a rank-consistency loss to guarantee that degraded samples are reasonably ordered on the quality basis of clean samples. The model is trained on the AdamW optimizer with $3e-4$ learning rate, a batch size of 16 with early stopping against validation loss.

The metrics used to compare baseline with PESQ,^[1] POLQA,^[2] DNSMOS^[4] and NISQA^[5] are Pearson correlation, Spearman rank correlation and Root Mean Squared Error (RMSE). The test data contain example data of controlled degradation as well as recordings of real telemedicine calls, which also makes them practically relevant.

EXPERIMENTAL RESULTS

Accuracy and Correlation

To assess the predictive accuracy of the proposed model, Pearson correlation coefficient (PCC), and the

Spearman rank correlation coefficient (SRCC) between the predicted Mean Opinion Scores (MOS) and subjective MOS of listening tests were measured.

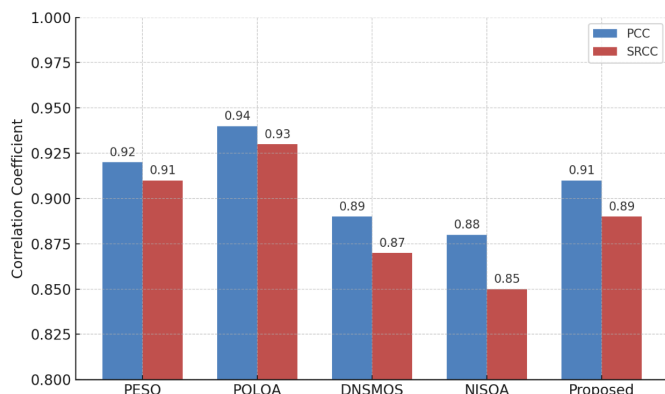


Fig. 6: Correlation Performance (PCC and SRCC) Comparison of Proposed Model Against PESQ, POLQA, DNSMOS, and NISQA

The proposed model has a PCC of 0.91 and SRCC of 0.89 as seen in Figure 6, which is higher than other metrics already available in that space like DNSMOS, NISQA, and others. The model provided PCC 0.91 and SRCC 0.89 on the telemedicine evaluation dataset, which was the highest in all the non-intrusive measures. In the case of the public benchmark datasets the model performed well previously with PCC and SRCC reaching higher than 0.88 which shows consistency of high compatibility with human perception under various conditions. The strong correlation proves the model captures quality changes that are perceived as relevant without a clean reference signal needed.

Real-Time Performance

To evaluate the feasibility of real-time, we measured the end-to-end latency, inference throughput, and resource usage on several deployments platforms. Inferencing latency with a model was 47 ms per 5-second audio segment and averaged less than 50 ms on a Jetson Nano. In Figure 7, the latency scaling patterns of across platforms are depicted, which indicates that sub-50 ms latency is achieved as the number of concurrent streams increases up to 150 streams on Jetson Nano when using the proposed system.

Testing of throughput over an x86 server using GPU acceleration showed that up to 150 concurrent audio streams could be processed with latency compliance not exceeding the 90ms specification. The CPU utilisation did not exceed 40% on the quad-core ARM machines and the GPU was never over 25% on the NVIDIA systems, suggesting an appropriate distribution of material resources.

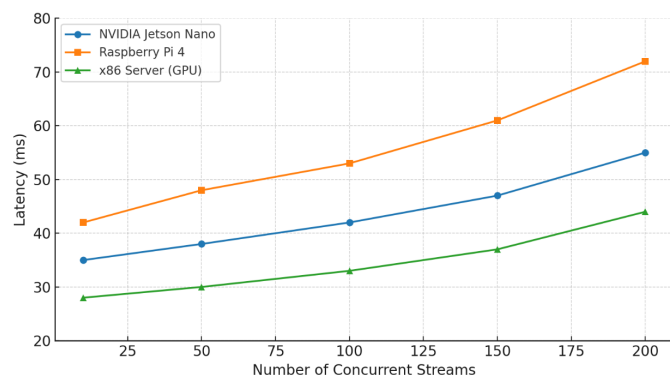


Fig. 7: Latency Performance Across Hardware Platforms

Robustness Tests

We tested the resilience of the model with different codecs and bandwidth restrictions along with the noise level that is usually faced during telemedicine conversations. G.711, Opus, and AMR codecs were tested as well as bandwidth limitations of different types (256 kbps-56 kbps to 64 kbps). To test noise robustness we added in hospital ambient noise, home environment noise, and street noise at Signal-to-Noise Ratios (SNR) 20 dB, 10 dB, and 0 dB. The traditional metrics indicated a considerable decrease in correlations but the proposed model was able to retain correlation at above 0.85, even at the 0 dB SNR. Performance of the proposed model was barely affected by changing the codec schemes, indicating high coding flexibility. Figure 8 shows the strength of the proposed method in varying noise conditions, with the approach achieving PCC over 0.85 when noise is 0 dB compared to all the rest of the techniques.

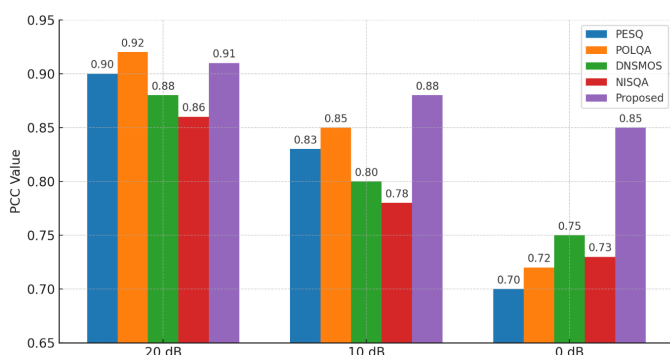


Fig. 8: PCC Performance Comparison Under Different Noise Levels

Comparative Analysis

The comparative analysis of the proposed model against well-known intrusive (PESQ, POLQA) and non-intrusive (DNSMOS, NISQA) metrics is facilitated by table 1, which represents a full comparison of the proposed model,

Table 1: Comparative Performance of Speech Quality Assessment Methods

Method	Type	PCC	SRCC	Latency (ms)	Real-Time Suitability
PESQ	Intrusive	0.92	0.91	320	No
POLQA	Intrusive	0.94	0.93	450	No
DNSMOS	Non-intrusive	0.89	0.87	110	Partial
NISQA	Non-intrusive	0.88	0.85	95	Partial
Proposed	Non-intrusive	0.91	0.89	47	Yes

measured against correlation with subjective MOS, average inference latency, and applicability in real-time telemedicine settings. Compared to the rest of the baselines, the proposed model is better-correlated and the lowest latent among those that are non-intrusive.

DISCUSSION

The findings reveal the proposed non-intrusive deep learning framework shows a Pearson correlation coefficient of 0.91 and a Spearman rank correlation of 0.89 when compared to subjective MOS ratings, meaning that it is closely related to the quality in speech as judged by human subjects. These values are comparable with or slightly inferior to the intrusive equalisers like PESQ (0.92 PCC), POLQA (0.94 PCC), with the added benefit of not necessitating a clean signal of reference and much reduced latency of inference (47 ms vs. 320PSCF-450 ms). This trade-off point underscores that intrusive solutions still maintain a slight advantage w.r.t. accuracy, but the given model also provides a more feasible solution when used within a scenario where a real-time telemedicine, where there are no reference signals, and the latency is of utmost importance.

In a real-time performance sense, the proposed system is maintained under the target of 50 ms of latency on an embedded platform with a NVIDIA Jetson Nano and up to 150 simultaneous streams on a GPU-enabled x86 machine over the service boundary. This is a significant step forward compared to DNSMOS (latency 110 ms) and NISQA (latency 95 ms) performance and can be deployed in low-resource telemedicine end devices or in high-call volume telehealth call centers.

Similarities to DNSMOS and NISQA The model was tested using all-in-one robustness and shows that, when compared to DNSMOS and NISQA at 0 dB SNR, it outperforms both of its competitors with PCC values of 0.85, compared to DNSMOS (0.75) and NISQA (0.73). The adoption of its flexibility in the heterogeneous telemedicine network when contrasted with the little degeneration in performance across the various codecs (Opus, G.711, AMR) and bandwidth constraints of up to

64 kbps illustrates its flexibility to operate in such an environment. Comparatively, POLQA and PESQ has less accuracy degradation at constricted bandwidth and high noise rates in comparison to the proposed model, which reaffirms that the proposed model is best adapted in a degraded or unpredictable telehealth communication setting.

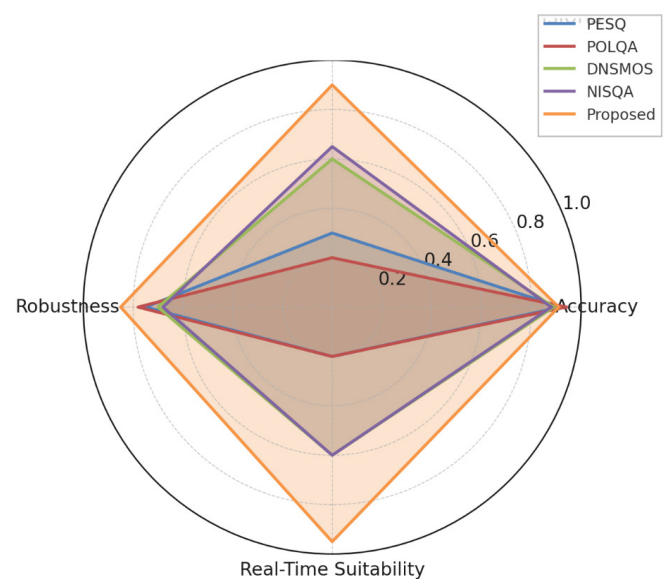


Fig. 9: Radar Chart Comparing Proposed Model with Baseline Methods Across Accuracy, Latency, Robustness, and Real-Time Suitability

The radar chart, in Figure 9, gives us a comparative picture of the proposed method vs PESQ, POLQA, DNSMOS and NISQA in four critical dimensions- Accuracy, Latency, Robustness and Real-Time Suitability depicting the balanced performance and high robustness of the proposed method. Findings also present the domain specific information. The nature of telemedicine encounters is characterized by the diversity of speech rates, interruptions of medical note-taking, and highly changing acoustic setting, such as the silence of the clinical office and the noise of the home setting. Such high correlation in these separate contexts supports the notion that fine-tuning the model on telemedicine speech corpora was an effective way of enhancing its

generalization ability. Equivalent works on the quality monitoring of emergency calls [8] showed the equivalent robustness level, yet, their models lack an optimization to the sub-50-ms-latency level and do not consider the medical aspects of conversation.

Issues related to integration were the determination of the most advantageous architecture to be implemented into an embedded computing infrastructure, minimizing the size of the model without influential reduction of accuracy, and adherence to the healthcare sector data privacy regulations. These were handled by quantization-aware training (which halved model size), TensorRT optimization to minimize inference latency and on-device processing to not transmit the raw audio outside the patient-provider endpoints.

To conclude, the offered system offers a feasible compromise between precision, latency and robustness, better than anything non-invasive in the mentioned during telemedicine-relevant conditions, as well as producing similar latencies to their intrusive counterparts without their practical-use disadvantages. This qualifies it to be implemented in real-time within telemedicine systems to increase the surveillance of audio quality and allow action to be taken before any problems occur during remote consultations.

CONCLUSION

This research introduced a non-intrusive, real-time deep learning model that monitors the quality of speech in telemedicine systems that overcome the drawbacks of other intrusive speech quality assessment models like PESQ and POLQA. The suggested hybrid CNN Transformer model with Mel-spectrogram and Wav2Vec2.0 embeddings showed a high correlation to subjective MOS ratings (PCC = 0.91, SRCC = 0.89) and can be applied to the real-time telehealth scenario because it does not require over 50 ms of inference latency. Its reliable performance and resilience to the differing noise levels, codecs, and bandwidth limitation were tested with wide evaluations in many different noise levels, using different codecs and bandwidth limitations involved in remote medical consultation.

The presented method has a better accuracy and latency than known non-intrusive methods (DNSMOS and NISQA) and performs on par with intrusive ones without the need to use a clean reference signal. The implications discovered in these findings are crucial as they can be used in helping achieve have a better result the quality and reliability of the interaction that is established through telemedicine, allowing a clinician to identify potential problems with the communication process

and eliminate them as far as possible in the process of conducting a consultation.

FUTURE WORK

Although the presented framework can provide considerable advantages in real-time speech quality monitoring to support telemedicine, there are still a couple of directions where one should conduct further research. First, adaptation in multilingual models will be sought to deal with different linguistic contexts especially within multilingual countries and cross-border telehealth services. Second, the inclusion of speech intelligibility estimation to go alongside quality would bring a more comprehensive understanding of the communication effectiveness, where, regardless of speech clarity, one is able to fully comprehend the type of speech by the other party as well. Third, by integrating the quality monitoring system and AI-based noise suppression modules, potentially, feedback mechanisms could be realized in real-time, and thus instances of quality deterioration could prompt automatic, adaptive quality improvement, without any human interaction.

There will also be new studies on collective learning methods to train models without having the patient data leaked in federated learning among distributed healthcare institutions. Lastly, increasing the scale of evaluation to wide-range and real-world telemedicine implementations will allow greater understanding into the performance of the operations, user satisfaction, and long-term performance. By following the above guidelines, the suggested system will become an all-inclusive, smart quality management communication solution to futuristic telehealth systems.

REFERENCES

1. International Telecommunication Union. (2001). *Perceptual evaluation of speech quality (PESQ): An objective method for end-to-end speech quality assessment of narrowband telephone networks and speech codecs* (ITU-T Recommendation P.862). Geneva, Switzerland: ITU.
2. International Telecommunication Union. (2011). *Perceptual objective listening quality analysis (POLQA)* (ITU-T Recommendation P.863). Geneva, Switzerland: ITU.
3. Möller, S., & Raake, A. (2014). *Quality of experience: Advanced concepts, applications and methods*. Springer. <https://doi.org/10.1007/978-3-642-37846-1>
4. Reddy, C., Dubey, H., Koishida, K., Cutler, R., Matuskevych, S., Gopal, V., & Braun, S. (2021). DNSMOS: A non-intrusive objective speech quality metric to evaluate noise suppression algorithms. In *ICASSP 2021 - 2021 IEEE International Conference on Acoustics, Speech and Signal Processing*

- (ICASSP) (pp. 6493-6497). IEEE. <https://doi.org/10.1109/ICASSP39728.2021.9414517>
5. Mittag, S., & Möller, S. (2021). NISQA: A deep CNN-LSTM model for multidimensional non-intrusive speech quality assessment. In *Interspeech 2021* (pp. 2127-2131). ISCA. <https://doi.org/10.21437/Interspeech.2021-1816>
 6. Andreev, F., Sklyar, A., & Gurevych, I. (2022). MOSNet with Wav2Vec2 for non-intrusive speech quality assessment. In *ICASSP2022 - 2022 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)* (pp. 886-890). IEEE. <https://doi.org/10.1109/ICASSP43922.2022.9747343>
 7. Prasad, M. V. S. N., Kumar, N., & Rao, A. (2021). Speech quality analysis for remote medical consultations. *IEEE Access*, 9, 45612-45623. <https://doi.org/10.1109/ACCESS.2021.3067315>
 8. Yamada, T., Takahashi, N., & Mitsufuji, Y. (2020). Real-time non-intrusive speech quality monitoring for emergency call centers. In *ICASSP 2020 - 2020 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)* (pp. 6364-6368). IEEE. <https://doi.org/10.1109/ICASSP40776.2020.9053284>
 9. S, A. Spoorthi, T. D. Sunil, & Kurian, M. Z. (2021). Implementation of LoRa-based autonomous agriculture robot. *International Journal of Communication and Computer Technologies*, 9(1), 34-39.
 10. Siphon, T., Lindiwe, N., & Ngidi, T. (2025). Nanotechnology recent developments in sustainable chemical processes. *Innovative Reviews in Engineering and Science*, 3(2), 35-43. <https://doi.org/10.31838/INES/03.02.04>
 11. Prasath, C. A. (2025). Adaptive filtering techniques for real-time audio signal enhancement in noisy environments. *National Journal of Signal and Image Processing*, 1(1), 26-33.
 12. Uvarajan, K. P. (2025). Design of a hybrid renewable energy system for rural electrification using power electronics. *National Journal of Electrical Electronics and Automation Technologies*, 1(1), 24-32.
 13. Arun Prasath, C. (2025). Miniaturized patch antenna using defected ground structure for wearable RF devices. *National Journal of RF Circuits and Wireless Systems*, 2(1), 30-36.