

Objective Evaluation of a Deep Learning-Based Noise Reduction Algorithm for Hearing Aids Under Diverse Fitting and Listening Conditions

Trends in Hearing

Volume 29: 1–26

© The Author(s) 2025

Article reuse guidelines:

sagepub.com/journals-permissions

DOI: 10.1177/23312165251396644

journals.sagepub.com/home/tia



Vahid Ashkanichenarlogh^{1,2} , Paula Folkeard^{1,3} , Susan Scollie^{1,3} , Volker Kühnel⁴ and Vijay Parsa^{1,2,3}

Abstract

This study evaluated a deep-neural-network denoising system using model-based design, comparing it with adaptive filtering and beamforming across various noise types, SNRs, and hearing-aid fittings. A KEMAR manikin fitted with five audiograms was recorded in reverberant and non-reverberant rooms, yielding 1,152 recordings. Speech intelligibility was estimated using the HASPI from 1,152 KEMAR manikin recordings. Effects of processing strategy and acoustic factors were tested with model-based within-device design that account for repeated recordings per device/program and fitting. Linear mixed model results showed that the DNN with beamforming outperformed conventional processing, with strongest gains at 0 and +5 dB SNR, moderate benefits at –5 dB in low reverberation, and none in medium reverberation. Across SNRs and noise types, the DNN combined with beamforming yielded the highest predicted intelligibility, with benefits attenuated under moderate reverberation. Azimuth effects varied; because estimates were derived from a better-ear metric on manikin recordings. Additionally, this paper reports comparisons using metrics of sound quality, for an intrusive metric (HASQI) and the pMOS non-intrusive metric. Results indicated that model type interacted with processing and acoustic factors. HASQI and pMOS scores increased with SNR and were moderately correlated ($r^2 \approx 0.479$), supporting the use of non-intrusive metrics for large-scale assessment. However, pMOS showed greater variability across hearing aid programs and environments, suggesting non-intrusive models capture processing effects differently than intrusive metrics. These findings highlight the promise and limits of non-intrusive evaluation while emphasizing the benefit of combining deep learning with beamforming to improve intelligibility and quality.

Keywords

hearing aids, deep neural networks, denoising, sound quality metrics, vent effects

Received: March 27, 2025; revised: September 6, 2025; accepted: October 28, 2025

Introduction

In our daily lives, we encounter sounds that help us perceive our surroundings and navigate through them. Some sounds, such as background noise, can also hinder speech clarity, increase listening effort, and reduce overall comfort (Krueger et al., 2017). For individuals with hearing loss, these challenges are further exacerbated, and even when using hearing aids, noisy environments can compromise speech comprehension and substantially increase cognitive load (Pichora-Fuller et al., 2016). To mitigate these difficulties, hearing aids employ signal processing strategies designed to attenuate unwanted noise and enhance desired speech signals. These strategies include microphone array

¹National Centre for Audiology, Western University, London, Canada

²Department of Electrical and Computer Engineering, Western University, London, Canada

³School of Communication Sciences and Disorders, Faculty of Health Sciences, Western University, London, Canada

⁴Sonova AG, Staefa, Switzerland

Corresponding author:

Vahid Ashkanichenarlogh, National Centre for Audiology, Western University, London, Canada; Department of Electrical and Computer Engineering, Western University, Elborn College, 1201 Western Rd, Room 2262, London, Canada.

Email: vashkani@uwo.ca



Creative Commons Non Commercial CC BY-NC: This article is distributed under the terms of the Creative Commons Attribution-NonCommercial 4.0 License (<https://creativecommons.org/licenses/by-nc/4.0/>) which permits non-commercial use, reproduction and distribution of the work without further permission provided the original work is attributed as specified on the SAGE and Open Access page (<https://us.sagepub.com/en-us/nam/open-access-at-sage>).

beamforming and/or single channel noise reduction techniques, discussed further below.

Single-channel noise reduction (NR) algorithms have long been a core component of hearing aid signal processing, aiming to improve the signal-to-noise ratio (SNR) by attenuating background noise without requiring spatial information. These algorithms typically operate in the time-frequency domain, where they apply gain reduction to or adaptive filtering of time-frequency regions dominated by noise (Boll, 1979; Scalart & Filho, 1996). Common approaches include spectral subtraction, Wiener filtering, and minimum mean-square error (MMSE) estimators. For a comprehensive review of single-channel NR in hearing aids see (Lakshmi et al., 2021).

Although NR algorithms are widely implemented in hearing aids with the goal of improving the listening experience, their impact on speech understanding remains inconsistent across studies and listening conditions (Keidser et al., 2007; Zaar et al., 2019). While these methods can reduce the annoyance of background noise (Brons et al., 2014) and have shown some benefit in lowering listening effort (Bentler & Chiou, 2006; Desjardins, 2016), they often provide limited or no improvement in speech intelligibility, particularly in non-stationary noise environments (Brons et al., 2014; Rhebergen et al., 2010). The main challenge lies in accurately distinguishing speech from noise, especially when both occupy overlapping frequency regions or fluctuate similarly over time. Thus, the traditional single-channel NR methods face several limitations, including slower response times (Scollie et al., 2016), poor performance at low SNRs (Ohlenforst et al., 2018), and a dependency on the characteristics of the input signal. As such, the adaptive filtering (traditional noise reduction) methods in hearing aids are generally understood to not significantly enhance speech intelligibility (Brons et al., 2014; Chong & Jenstad, 2018; Völker et al., 2015).

One of the most effective strategies to enhance speech perception for hearing aid users is the use of beamforming processing. This technique involves utilizing the spatial characteristics of incoming sound to improve the SNR, thereby facilitating better access to speech signals in complex acoustic scenes. In hearing aids, beamforming processing leverages microphone arrays to enhance sounds arriving from a specific direction, typically the front, while suppressing noise from other directions. Key components of beamforming processing systems in hearing aids include fixed or adaptive beamforming microphones, gain control strategies that attenuate background noise, and complementary digital noise reduction algorithms aimed at improving the SNR (see reviews by Launer et al., 2016 and Lansbergen & Dreschler, 2022). In early digital beamforming hearing aids, this approach was shown to improve the clarity of frontally presented speech by attenuating noise from lateral and rear directions (Appleton & König, 2014; Boymans & Dreschler, 2000; Picou et al., 2014). Although beamforming can significantly improve speech intelligibility in noisy environments, it may also attenuate desired sounds coming from off-axis directions or alter binaural

cues, particularly with monaural beamforming. As a result, users may experience reduced spatial awareness and difficulty following conversations in dynamic or multi-speaker scenarios. More recent developments to address these deficiencies include adaptive, binaurally-linked, asymmetrical, and multi-beam beamforming systems that offer increased flexibility and spatial selectivity in dynamic acoustic environments (Andersson et al., 2021; Folkeard et al., 2024).

In summary, common hearing aid noise reduction strategies involve gain reduction and/or adaptive filtering, which dynamically adjusts filter parameters to suppress time-varying noise, and spatial filtering through beamforming, which enhances sound from desired directions while suppressing unwanted signals from non-target areas (Kollmeier & Kiessling, 2018; Lansbergen & Dreschler, 2022). These processors are widely implemented and demonstrate effectiveness primarily in limited scenarios, such as with steady-state background noise or when the speech source is spatially separated and located in front of the listener as is the case in many controlled laboratory settings (Kumar et al., 2023). Their effectiveness diminishes substantially in diffuse or highly dynamic environments where spatial separation is less distinct. As such, studies, e.g., (Picou, 2022), continue to highlight the real-world challenges faced by hearing aid users, emphasizing the need for effective noise reduction and beamforming to improve speech intelligibility and listening comfort in diverse acoustic environments.

Recent progress in deep learning has impacted numerous areas (Abdel-Jaber et al., 2022; Sharifani & Amini, 2023). Within the context of hearing aids, the advent of deep learning has led to acoustic scene classification, which is a common feature that helps to automate multiprogram use in hearing aids (Jiang et al., 2020; Ye et al., 2018). More recent applications using DNN based models have been introduced in fields such as speech recognition and enhancement (Ashkanichenarlogh & Parsa, 2024; Diehl et al., 2022; Rojc & Mlakar, 2022).

Earlier studies have shown that deep learning techniques for denoising (Goehring et al., 2016; Zhao et al., 2018) and separating overlapping speakers (Bramsløw et al., 2018) can improve speech clarity for users of cochlear implants (Goehring et al., 2017) and those with severe to profound hearing loss under constant SNR. Similarly, Bentsen et al. (2018) showed early promise by combining deep neural networks with ideal ratio mask estimation in computational speech segregation. Their results suggested significant improvements in intelligibility for listeners, validating the integration of deep learning with classical masking-based approaches. Recent research does suggest that deep learning-based models hold considerable promise for reducing background noise and substantially improving speech clarity for hearing aid users (Andersen et al., 2021; Diehl et al., 2023). Numerous recent studies have demonstrated that DNN-based algorithms can enhance speech intelligibility for hearing-impaired listeners, including those using cochlear implants, even in highly challenging environments such as multi-talker babble and reverberant conditions.

Keshavarzi et al. (2019) directly compared a Recursive Neural Network (RNN)-based noise reduction algorithm with spectral subtraction (SS) using both subjective ratings and objective intelligibility measures. Mild-to-moderate hearing-impaired listeners consistently preferred the DNN approach over SS and unprocessed speech in terms of both intelligibility and quality. Interestingly, SS did not offer significant benefit over unprocessed speech, highlighting the limitations of traditional methods in babble noise. Borjigin et al. (2024) also showed that an RNN-based model significantly improved speech intelligibility for cochlear implants users in both stationary and non-stationary noise conditions, settings where conventional methods typically fail. Similarly, Gaultier and Goehring (2024) evaluated DNN-based enhancement using multiple microphones and demonstrated significant improvements in speech reception thresholds of up to 10.3 dB for cochlear implants listeners in noisy-reverberant conditions, with greater benefits observed when spatial cues were utilized.

Healy, Taherian et al. (2021) specifically addressed the feasibility of causal DNN models for real-time applications. Their Dense-UNet and temporal convolutional network achieved a mean benefit of 46.4 percentage points for hearing-impaired listeners in reverberant, talker-interference scenarios, even after conversion to causal operation. The loss of performance due to causality was statistically significant in only the most complex conditions, underscoring the real-world applicability of these models. In another study, Healy et al. (2023) documented substantial progress in the efficacy and generalizability of DNN-based noise reduction over the past decade. By comparing a recurrent neural network with traditional noise reduction methods, they showed that even under causal (real-time) constraints and mismatched training/testing conditions, the algorithm still delivered intelligibility gains averaging 51 percentage points for hearing-impaired listeners.

In the broader context of monaural speech enhancement, Zheng et al. (2023) provided a comprehensive review of sixty years of research, emphasizing the leap in performance and robustness introduced by deep learning methods. Their work illustrates how deep learning has transformed frequency-domain approaches, enabling them to surpass traditional algorithms in most relevant metrics. Thoidis and Goehring (2024) further demonstrated that DNNs can enhance the intelligibility of a target speaker in noisy multi-talker environments for both normal-hearing and hearing-impaired listeners. Their study supports the generalization of DNN enhancement benefits across different listener groups and listening conditions.

DNN-based models have required adaptation to work within the limited computational power available during real-time use on mobile and portable devices like hearing aids. Recently, noise reduction methods for HAs have been developed using DNN-based algorithms (Andersen et al., 2021; Schröter et al., 2022). For example, Christensen et al. (2024) compared real-life listening experiences of experienced hearing aid users wearing premium hearing aids with

traditional noise reduction and with DNN-based noise reduction. While overall satisfaction did not differ, the DNN device provided more consistent sound satisfaction regardless of background noise, whereas the satisfaction ratings for the non-DNN device depended on ambient SNR.

Given that deep learning systems tend to scale effectively with an increase in parameters and computational power (Tay et al., 2022), downsizing them often leads to a decline in output quality. As hearing aids evolve, the integration of advanced, on-device denoising systems is becoming an increasingly realistic goal. Evaluating the performance of current first-generation denoising technologies is therefore a critical step, not only to assess their immediate impact on speech quality and listening comfort, but also to inform the design and optimization of future, more powerful systems. Early insights from these evaluations will help shape how next-generation models are developed, implemented, and ultimately integrated into everyday hearing aid use. Numerous computational metrics have been developed to evaluate speech and sound quality, with well-established examples including the perceptual evaluation of speech quality (PESQ) (Falk et al., 2015) and the short-time objective intelligibility (STOI) (Taal et al., 2011). However, these metrics fall short in accurately evaluating signals processed by hearing aids because they do not account for the specific impact of hearing loss. To address this limitation, metrics that include hearing loss models have been developed, such as the Hearing-Aid Speech Quality Index (HASQI) (Kates & Arehart, 2014) and the Hearing-Aid Speech Perception Index (HASPI) (Kates & Arehart, 2021). HASQI is designed to measure speech quality, while HASPI focuses on assessing speech intelligibility. A significant limitation of most of these metrics is that they are intrusive, meaning they require a clean reference audio signal to compare against the processed audio, in order to estimate the differences and assess the impact of processing. This reliance on a reference signal can pose challenges when a clean reference may not be available. Non-intrusive metrics such as the pMOS (Diehl et al., 2022) may provide alternatives for sound quality estimation.

This study evaluates a DNN-based noise reduction algorithm implemented directly on Phonak Audéo I90-Sphere receiver-in-the-canal hearing aids, building on prior work (Diehl et al., 2022, 2023). All signal processing strategies were tested on the same hearing aid hardware platform to ensure consistency. The DNN-based algorithm was benchmarked against standard methods, including a single-channel Wiener filter algorithm, a binaural beamforming strategy, and combinations thereof. A total of six processing conditions were tested across three SNR levels (−5, 0, and 5 dB) and two types of background noise: stationary speech-shaped noise and multi-talker babble. Hearing-aid gains were fitted to four standard audiograms (N2–N5) using Adaptive Phonak Digital 3.0 with 100% prescribed gain. Both occluded and vented conditions were evaluated. The novel DNN-based algorithm investigated in this study may provide advantages

by being less dependent on the signal type, and by functioning effectively regardless of the noise source's location, compared to adaptive and beamforming filters. This approach has the potential to enhance hearing aid performance in a broader range of acoustic environments.

Methodology

Purpose

The primary aim of this study was to evaluate the performance of a new DNN-based denoising algorithm for hearing aids across different SNRs in listening environments with low and medium reverberation, using objective metrics. The evaluation considered the secondary effects of venting configurations, types of background noise, and signal azimuth. Comparisons were made with adaptive filtering and beamforming-based noise reduction methods. Objective metrics were used to estimate both speech intelligibility and sound quality, with a secondary aim of assessing the effectiveness of intrusive and non-intrusive measures for evaluating hearing aid denoising and/or beamforming algorithms. We used the TRIPOD checklist when writing this report (Collins et al., 2015).

Methods

In this study we have conducted experiments in three SNR conditions (-5 , 0 and 5 dB SNR), and two background noise types: speech-shaped noise (SSN) and multi-talker babble (MTB) noise, where the SSN was a stationary noise, spectrally shaped to mimic the shape of human speech, making it particularly effective at obscuring speech when the background contains similar speech elements (Renz et al., 2018). The MTB (Krishnamurthy & Hansen, 2009) was a six-talker babble where the condition was specifically designed to include three unique sets of 2-talker babble, with each pair of talkers consisting of one male and one female and described below.

The SNR was measured separately in both the reverberant room and the sound booth to account for the distinct acoustic characteristics of each environment. In both cases, the measurements were conducted with a KEMAR manikin positioned at the center of the room, and surrounded by four loudspeakers placed at fixed angles. Prior to the SNR measurements, a thorough calibration procedure was carried out to ensure consistent output levels across all loudspeakers. To determine the SNR, a speech signal was first presented through the loudspeakers, and its level was adjusted to achieve a calibrated sound pressure level (SPL) of 70 dB at the position of the manikin's ear, measured using SpectraPLUS software. Following this, a multi-talker babble or equivalent noise was introduced, and the SPL of the noise was adjusted to 65 dB(A) at the same measurement point. The SNR was then calculated as the difference between the speech

and noise levels, resulting in an SNR of $+5$ dB. This process was repeated for both acoustic environments, ensuring consistency in setup and measurement methodology.

In our experiments, Phonak Audéo I90-Sphere I90 receiver-in-the-canal hearing aids were used. Hearing-aid gains were fitted using the Adaptive Phonak Digital 3.0 fitting formula and set to 100% gain to standard audiograms N2-N5 (Bisgaard et al., 2010). This study evaluated six distinct hearing aid signal processing strategies, with an emphasis on understanding the performance of a new DNN-based denoising algorithm in comparison with conventional noise reduction and beamforming microphone techniques. The DNN-based speech denoising system designed to separate speech from background noise in real time is previously described by Diehl et al. (2023). Its commercial implementation within the tested hearing aids is described by Hasemann and Krylova (2024). Here, the DNN was implemented on a dedicated 50 MHz co-processor chip that is activated within a program dedicated to processing speech in loud noise. The denoising network has a U-Net architecture and predicts a complex-valued ideal ratio mask in 64 frequency bins from the short-time Fourier transform of the noisy speech signals. The extracted speech signal output of the co-processor chip then is fed to standard audio processing path.

The tested signal processing strategies are summarized in Table 1 and described in more detail below.

No Processing (Omnidirectional): This baseline condition involved no active noise reduction or beamforming microphone modes. The microphone mode was set to omnidirectional, providing a baseline for comparing the effects of other algorithms.

NoiseBlock (Omnidirectional): This condition used Phonak's default implementation of NoiseBlock, a single-channel Wiener filter-based statistical noise reduction algorithm. It operates by estimating the noise spectrum during speech pauses and suppressing noise while preserving speech, with the microphone mode maintained in omnidirectional mode.

Beam (StereoZoom Beamforming): This condition activated StereoZoom, a binaural adaptive beamforming algorithm that uses information from both hearing aids to enhance the signal from the front (0° azimuth) and suppress sounds from other directions. NoiseBlock was disabled in this condition.

Beam + NoiseBlock: In this condition, StereoZoom beamforming was combined with NoiseBlock statistical noise reduction. This configuration reflects common clinical use where beamforming processing and noise reduction are both engaged.

DNN (Omnidirectional): This condition applied a proprietary deep-neural-network-based denoising algorithm that operates in parallel to other signal processing within this hearing aid. The DNN model was implemented within the hearing aids and operated in real-time. The microphone mode remained omnidirectional to isolate the effects of the DNN algorithm from those of beamforming microphones.

Table 1. Hearing Aid Signal Processing Conditions. All Signal Processing Conditions Were Tested at Their Default Settings for N2, N3, N4, and N5 Audiograms.

Signal Processing	Description
No Processing	No noise reduction method has been applied. Microphone mode: Omni
NoiseBlock	A statistical noise-canceling method has been applied. Microphone mode: Omni
Beam	Adaptive binaural beamforming (Stereo Zoom) has been applied.
Beam + NoiseBlock	Statistical noise canceling with adaptive binaural beamforming (Stereo Zoom) has been applied.
DNN	DNN denoising has been applied. Microphone mode: Omni
Beam + DNN	DNN denoising has been applied. Microphone mode: Monaural beamforming

Beam + DNN (Monaural Beamforming Microphone):

In this condition, the DNN denoiser operated in series on the beamforming microphone output. The microphone mode was adjusted to a monaural beamforming mode with fixed frontal directionality. This condition aimed to examine the impact of combining machine learning-based denoising with beamforming microphone processing versus other signal processing strategies.

Venting: All audiograms were tested with earmolds customized for the KEMAR manikin's ear to create a non-vented control condition. Additionally, a vented fitting was created for each of the audiograms. Vent configurations were adjusted to simulate clinical vent assignments (N2: open dome; N3: vented dome; N4 and N5: 1 mm vent). Details of the audiograms and venting conditions are displayed in Table 2.

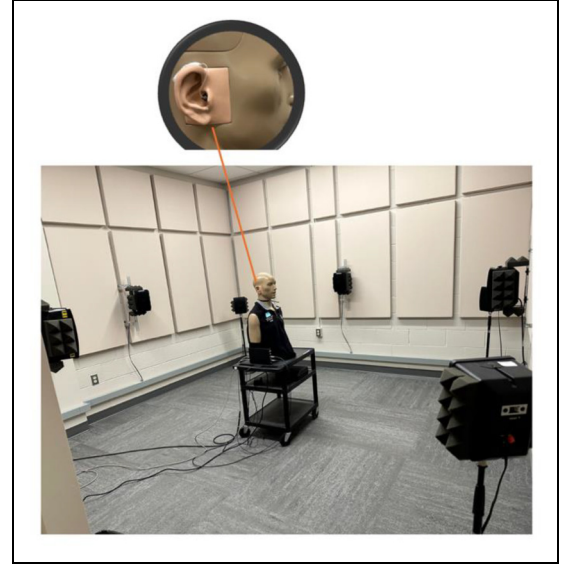
To evaluate the performance of these hearing aids, they were positioned on a KEMAR manikin equipped with the G.R.A.S. KB0065 pinna for the right ear and the G.R.A.S. KB0066 pinna for the left ear, both of which are designed to represent large ear geometries in accordance with standardized human ear models. The KEMAR was surrounded by four speakers positioned at 0°, 90°, 180°, and 270°.

Azimuth of speech: Two different talker-azimuth configurations were employed: one with speech presented from the front (0°) and the other with speech presented from the side (90°). Background noise was emitted from the other three speaker locations (i.e., Configuration 1 – Speech at 0°, Noise from 90°, 180°, 270°, Configuration 2 – Speech at 90°, Noise from 180°, 270°, 0°). Figure 1 illustrates the setup of our experiments in the reverberant room.

SNR: Recordings were captured with speech presented at 65 dB SPL and SNRs set at 5 dB, 0 dB, and -5 dB. This range of SNRs allowed evaluation of the hearing aids' performance under varying acoustic conditions representing difficult listening situations. These SNR values align with typical noisy real-life situations faced by hearing aid users (Lelic et al., 2022; Weisser & Buchholz, 2019; Wu et al., 2018).

Table 2. Audiograms and Venting Conditions.

Audiograms	Receiver	Vented	Closed
N2	M	Open Dome	Kemar Mold - no vent
N3	M	Vented Dome	Kemar Mold - no vent
N4	M	Kemar Mold 1 mm	Kemar Mold - no vent
N5	P	Kemar Mold 1 mm	Kemar Mold - no vent

**Figure 1.** Experimental Setups the KEMAR Manikin with Phonak Audéo I90-Sphere I90 Hearing Aids in a Room with Moderate Reverberation.

The 5 dB SNR conditions corresponds to moderately noisy settings, such as conversations in quieter public spaces where background noise is present but speech remains somewhat dominant, while 0 dB reflects highly challenging scenarios, such as crowded areas or social gatherings, where the speech and noise levels are comparable, and -5 dB represents difficult listening conditions, such as heavy traffic or noisy restaurants, where the background noise is louder than the target speech, posing significant difficulties for understanding speech.

Stimuli: The speech stimuli used in the experiments consisted of passages from the Connected Speech Test (CST: Saleh et al., 2020), presented either in SSN or in MTB emanating from the speaker locations described above. The babble track was created from these talkers reading unused CST passages on the topics of "Umbrella," "Vegetable," "Camel," "Glue," and "Cactus" (Saleh et al., 2020).

Acoustic environments and reverberation time: The experiments were conducted in two distinct acoustic environments: an audiometric double-walled sound booth with minimal reverberation and a room with moderate reverberation. The reverberant room, depicted in Figure 1, had a measured reverberation time (T30) of 400 ms. These environments were chosen to evaluate the DNN-based noise reduction performance under varying acoustic conditions,

with a particular focus on the effects of reverberation, signal azimuth, and acoustic coupling.

Metric computation and statistical analyses: For each test condition, segments pertaining to individual sentences from the CST passage recordings from the left and right hearing aids were extracted. The HASPI and HASQI scores for each of these extracted segments were calculated using the procedures outlined by (Kates & Arehart, 2022) averaged across all 9 sentences from each speech sample. The pMOS score was computed by uploading the extracted segments to the metric website (<https://metric.audatic.ai/>, Diehl et al., 2022) and retrieving the cloud-computed pMOS value. Metric values were chosen as the highest value obtained from either ear, referred to as “better-ear” values henceforth. The HASPI, HASQI, and pMOS values averaged across the sentence segments were used for statistical analyses.

To assess the performance of the DNN-based noise reduction algorithm in comparison to alternative hearing aid signal processing strategies, we conducted both primary and secondary statistical analyses on the metric data. The primary analysis focused on examining the effects of processing program and SNR within each listening environment. The secondary analyses explored the influence of noise type, sound source azimuth, and venting condition across the same program and SNR combinations. All data were analyzed using Linear Mixed Models (LMM) implemented in IBM SPSS Statistics (v29.0). Fixed effects included program, SNR, noise type, azimuth, and venting, while random effects accounted for repeated measures within audiograms. Model diagnostics included inspection of QQ plots of the residuals and scatter plots of predicted values versus residuals (please see the supplemental data in the repository link). The QQ plots indicated approximate normality of residuals. The effects of heteroscedasticity were checked using cluster-robust standard errors (clustered by audiograms). Bonferroni corrections were applied to control for multiple comparisons, with an experiment-wise type 1 error rate of 0.05. Furthermore, in paired comparisons results discussed in the following section, differences in HASPI values between two processing conditions were reported only when they reached statistical significance ($p < 0.05$).

Results

As described above, statistical analyses were conducted independently for the HASPI, HASQI, and pMOS outcome measures. Among these, the HASPI data demonstrated greater sensitivity across test conditions and program settings, yielding a higher number of statistically significant main effects and interactions. This finding aligns with expectations, as HASPI values tend to exhibit greater variability across processing conditions within the range of SNRs examined in this study (see Kates & Arehart, 2022, for representative data). Consequently, the Results section emphasizes the analysis and interpretation of HASPI outcomes, while findings

related to HASQI and pMOS are presented in a more concise manner. The interested reader is referred to the supplemental data in the repository link for statistical analyses outputs for HASQI and pMOS data.

Effects of Hearing Aid Processing across SNR, per Listening Environment:

Low-Reverberation: The LMM was fitted for the HASPI scores for testing completed in the low-reverberation listening environment (Figure 2 (a)). It is noticeable that in all figures, each bar represents the mean and standard deviation across $n = 9$ sentences of each speech sample. Program and SNR were considered the Fixed Effects, and Audiogram was a random effect. Results revealed significant main effects of SNR and program type ($p < .05$), along with a statistically significant interaction between these factors ($F(10, 555) = 4.175$, $p < .001$). Under the most adverse SNR condition, beamforming plus DNN program significantly outperformed non-DNN processing conditions, yielding HASPI improvements of 11–13%, but did not differ from the DNN-alone condition (Supplemental Sound Booth Output). At 0 dB SNR, the beamforming plus DNN configuration outperformed all other programs, including DNN alone, by 18–33%, while DNN alone exceeded non-beamforming/non-DNN programs by 12%. At +5 dB SNR, the pattern was similar: beamforming plus DNN outperformed all but traditional beamforming, with gains of 16–30%, and DNN alone outperformed non-beamforming conditions by 20%. Traditional beamforming programs (Beam and Beam + NoiseBlock) also outperformed baseline at 0 and +5 dB SNR, whereas NoiseBlock alone showed no significant advantage at any SNR.

Medium-Reverberation: The LMM was fitted for the HASPI scores for testing completed in the medium-reverberation listening environment (Figure 2 (b)). Program and SNR were treated as fixed effects, with audiogram included as a random effect. Results revealed significant main effects of both SNR and program type ($p < .05$), along with a significant two-way interaction between SNR and hearing aid program type ($F(10, 555) = 2.219$, $p = .016$). In contrast to the findings in the low-reverberation room, post-hoc analyses for the –5 dB SNR condition showed no significant differences between any of the signal-processing schemes ($p > .05$), including comparisons between DNN alone, beamforming plus DNN, and the other programs (Supplemental Reverb Room Output). At the 0 dB SNR condition, the beamforming plus DNN program achieved the highest HASPI scores, outperforming all other processing conditions by 23–25%. The DNN-alone program also exceeded all non-DNN configurations, including both traditional beamforming and noise reduction programs, by 14–17%. At the +5 dB SNR condition, the beamforming plus DNN program achieved the highest HASPI scores, outperforming all non-DNN configurations by 16–18%. The DNN-alone program also exceeded all non-DNN conditions by 11–14%. In contrast to the low-reverberation results, none of the traditional noise-reduction or beamforming programs (NoiseBlock, Beam, or Beam +

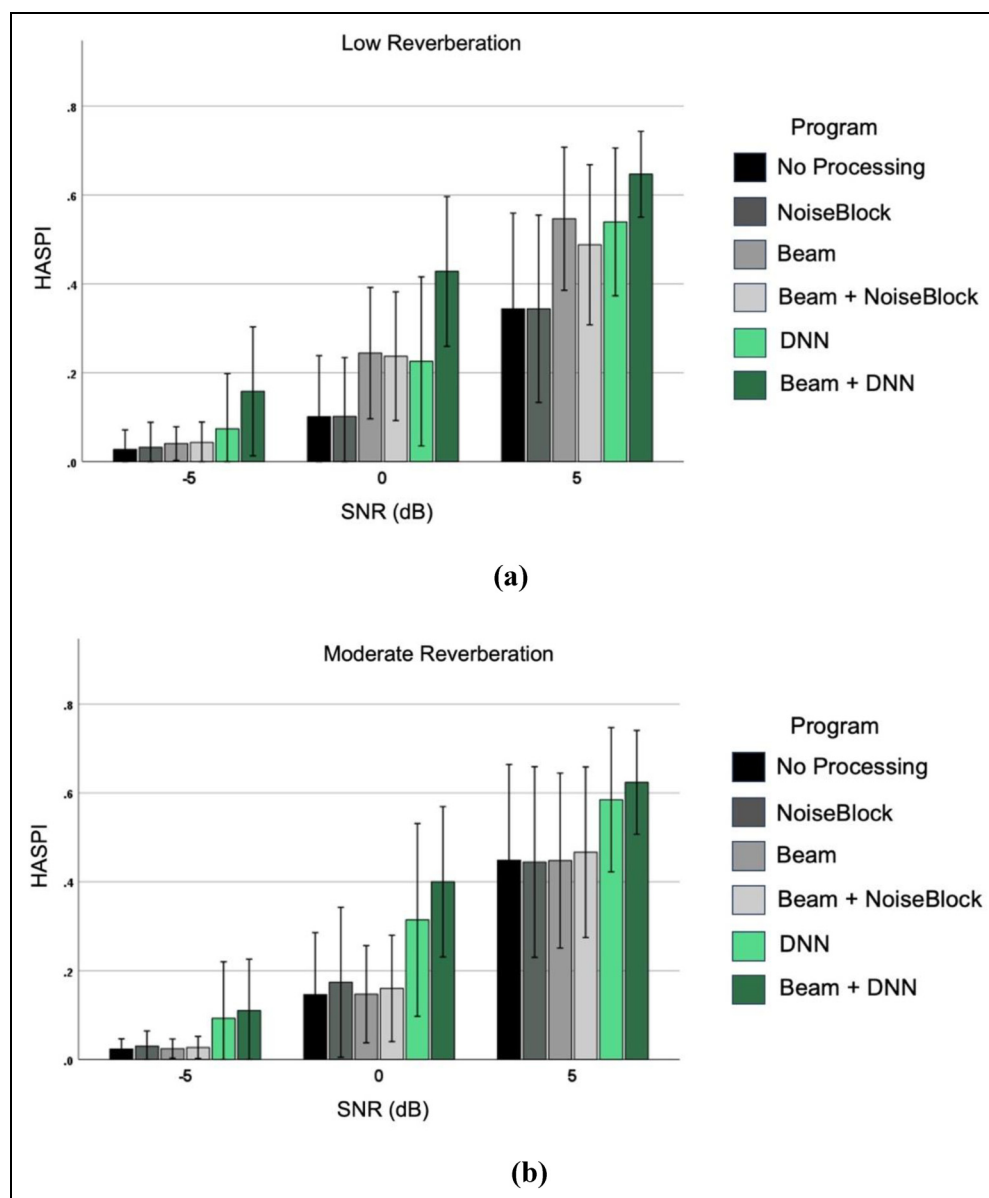


Figure 2. Effects of Hearing Aid Processing across SNR per Listening Environment. (a) Low Reverberation, and (b) Moderate Reverberation. Each Bar Represents the Better-Ear Average HASPI Score Calculated Across 9 Sentences of Each CST Passage. Error Bars Represent Standard Deviations Across Sentences.

NoiseBlock) showed significant improvements over the no-processing baseline at any SNR.

Effects of Hearing Aid Processing across Vent Types and SNR:

The LMM was fitted for the HASPI scores for testing completed in both listening environments, by SNR and Vent Type (see Figure 3). Program, SNR, and Vent Type were considered as Fixed Effects, and Audiogram was a random effect. As expected, the main effects of SNR and Vent Type were significant ($p < 0.05$), with lower HASPI scores for lower SNRs or more open hearing aid styles (Supplemental Venting Output). These two factors also interacted $F(6,1076.99) =$

9.94, $p < 0.001$), with HASPI scores that were not only lower but also more similar across vent types at the lowest SNR; this result is not unexpected (Wolfe et al., 2009) nor is it specific to the aims of this study. Results revealed a significant main effect for the program setting ($p < 0.05$), as well as a two-way interaction between venting and hearing aid program type ($F(15, 1077) = 3.47$, $p < .001$).

In the occluded condition, the beamforming plus DNN outperformed the non-DNN conditions by 20–28%, and the DNN only condition by 11%. The DNN only program outperformed all non-DNN programs by 9–16%. In general, the programs with beamforming had smaller mean differences from

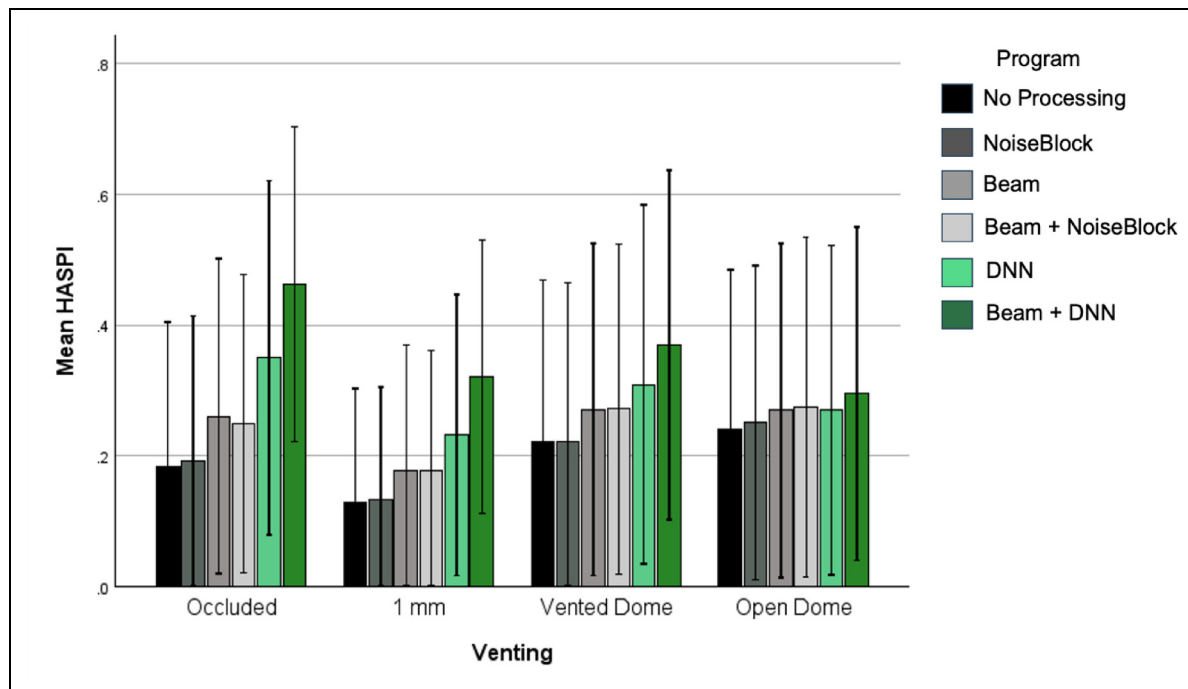


Figure 3. Effects of Hearing Aid Processing Across Different Vent Types. Each Bar Represents the Better-Ear Average HASPI Score Calculated Across Sentences of the CST Passage for That Test Condition. Error Bars Represent the Corresponding Standard Deviations.

the beamforming plus DNN, compared to those that did not. For vented conditions, small differences between the DNN and non-DNN programs were observed. For example, at the 1 mm vented condition, the beamforming plus DNN outperformed the non-DNN programs by a similar magnitude (14–19%). For the vented dome conditions, the beamforming plus DNN was significantly better than No Processing and NoiseBlock programs, and statistically similar to others. With an open dome fitting, none of the differences were statistically significant. Across conditions, the beamforming plus DNN was able to outperform a No Processing condition for the occluded, 1 mm, and vented dome fittings. In contrast, the traditional beamforming and noise reduction program settings outperformed the no processing condition only when fully occluded. The DNN only program performed in between the two conditions, outperforming the no processing condition when occluded and with 1 mm of venting, but not for the other vent conditions.

Effects of Hearing Aid Processing across Noise Types and SNR:

The LMM was fitted for the HASPI scores for testing completed in both listening environments, by SNR and Noise Type (Figure 4), Program, SNR, and noise type were considered as fixed effects, and audiogram was a random effect. As expected, the main effects of SNR and noise type were significant ($p < 0.05$), with lower HASPI scores for lower SNRs or more modulated background noises (Supplemental Noise Type Output). These two factors also interacted ($F(2,1113.0) = 5.48$, $p = 0.004$). In addition, there was a

significant main effect of the hearing aid program setting ($p < 0.05$), and a two-way interaction between noise type and hearing aid program setting ($F(5, 1113) = 5.78$, $p < .001$). Performance varied depending on the type of noise and processing strategy. Specifically, adaptive noise filtering, DNN alone, and DNN combined with beamforming all showed reduced effectiveness in MTB compared to SSN. The traditional beamforming settings (Beam and Beam + NoiseBlock) did not differ from the No Processing setting in the SSN condition but were significantly better than No Processing in MTB environments. The highest performance was observed with DNN and beamforming plus DNN in the SSN condition. These programs differed by about 30% with the most improvement for beamforming plus DNN, compared to the No Processing condition.

Effects of Hearing Aid Processing across Azimuth and SNR:

The LMM was fitted for the HASPI scores for testing completed in both listening environments, by SNR and Azimuth (Figure 5). Program, SNR, and Azimuth were considered as fixed effects, and audiogram was a random effect. The main effects of Program, SNR and Azimuth were significant ($p < 0.05$), and all had two-way interactions as well (Supplemental Azimuth Output). These three factors interacted significantly ($F(10, 1113) = 6.341$, $p < .001$). At the lowest SNR, the beamforming plus DNN outperformed the non-beamforming programs by about 10%, and there were no significant differences between the other conditions. At a 90-degree azimuth, the difference increased, with significant

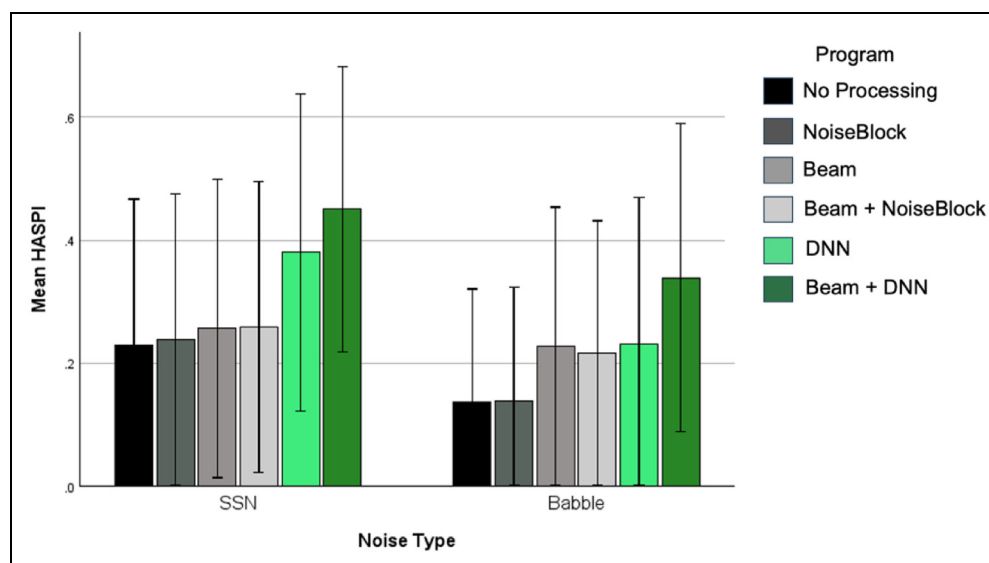


Figure 4. Effects of Hearing Aid Processing Across Noise Types and SNR. Each Bar Represents the Better-Ear Average HASPI Score Calculated Across Sentences of the CST Passage for That Test Condition. Error Bars Represent the Corresponding Standard Deviations.

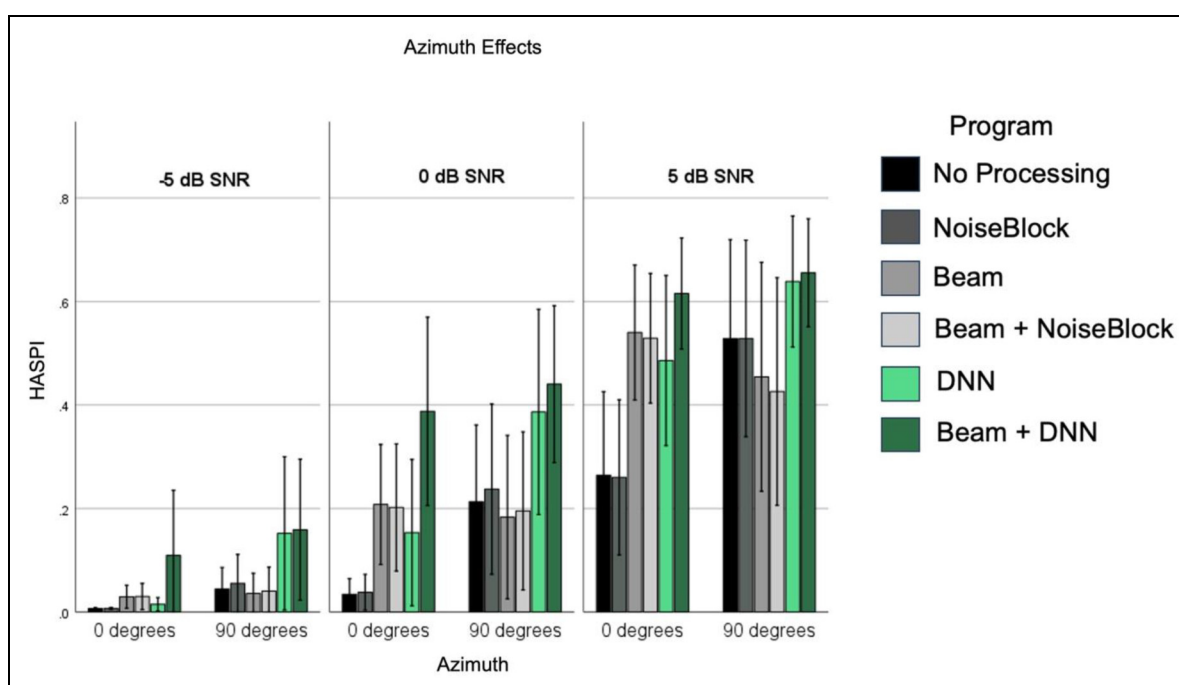


Figure 5. Effects of Hearing Aid Processing Across Azimuth and SNR. Each Bar Represents the Better-Ear Average HASPI Score Calculated Across Sentences of the CST Passage for That Test Condition. Error Bars Represent the Corresponding Standard Deviations. “Note: HASPI Values Reflect a Better-Ear Metric; Lateral (90°) Advantages may be Amplified by Head-Shadow and Should Not Be Taken to Imply Superior Bilateral Performance”.

beamforming plus DNN outperformance over the non-DNN programs by about 10–12%. DNN alone also outperformed the non-beamforming programs when the signal was at 90 degrees. At higher SNRs, more of the processing-based programs outperformed the no-processing condition, and the magnitude of the DNN advantage increased as follows. The

beamforming plus DNN was 35% better than the non-beamforming programs at 0 degrees, and at 90 degrees DNN programs outperformed all other programs by more than 20% (with beamforming) and by about 15–20% (without beamforming) at 90 degrees. Similar benefits with higher overall HASPI scores were observed at the +5 dB SNR. At

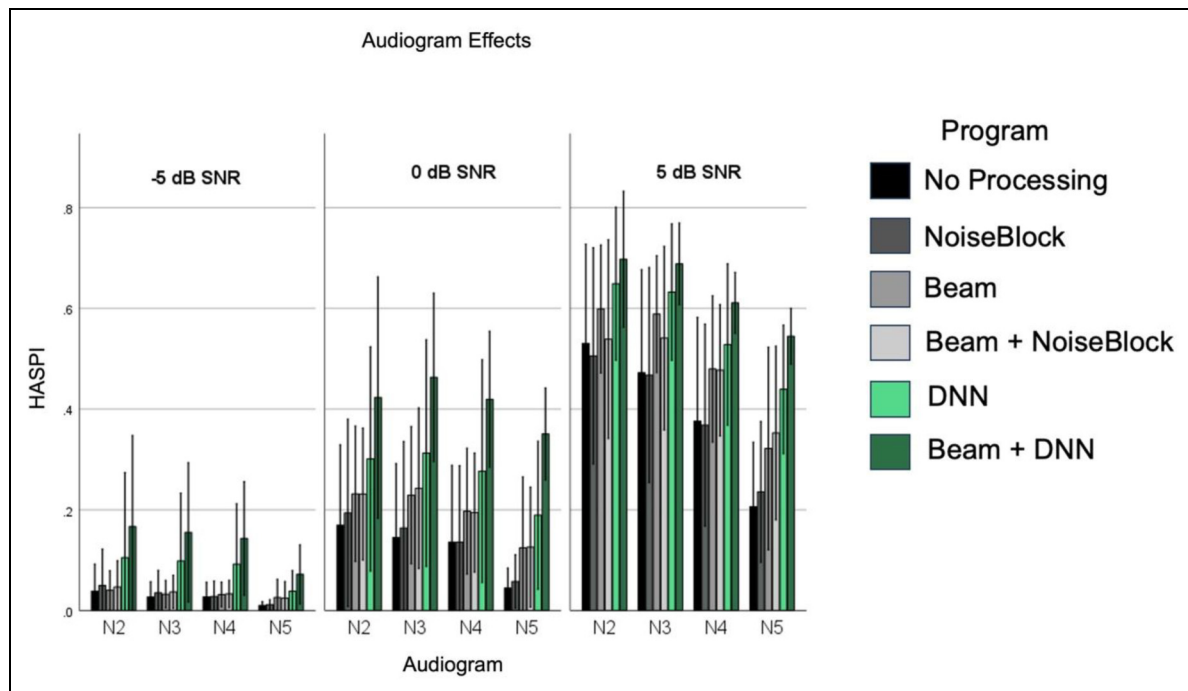


Figure 6. Effects of Hearing Aid Processing Across Audiograms and SNR. Each Bar Represents the Average HASPI Score Calculated Across $n = 9$ Sentences of the CST Passage for that Test Condition. Error Bars Represent the Corresponding Standard Deviations.

the zero degrees condition, the benefit for DNN was observed when it was combined with the beamformer, and poorer performance when used with omnidirectional processing. Figure 5 also highlights the performance of program settings involving traditional beamforming. With speech incidence at 0 degrees, the Beam and Beam + NoiseBlock settings were significantly better than the No Processing setting at 0 dB and 5 dB SNRs. At the same SNRs, however, their performance was degraded compared to the No Processing setting when speech emanated from the 90 degrees azimuth. These results are in line with previous studies [e.g., Hornsby and Ricketts (2007)] where the front-facing beamformer resulted in reduced recognition scores for speech stimuli in the presence of background noise, when the speech was presented from 90 degrees azimuth.

Effects of Hearing Aid Processing across Audiograms and SNR:

Figure 6 illustrates the impact of various hearing aid processing programs on HASPI scores across three SNR conditions (−5 dB, 0 dB, and 5 dB) and different audiograms (N2–N5). Consistent with analyses already presented above, performance is dependent upon SNR, and the DNN provided higher HASPI scores both overall and at lower SNRs compared to traditional processing. This study did not have an adequate sample of simulated listeners to perform a statistical analysis of the effects of audiogram. However, the descriptive data shown in Figure 6 illustrates a decline in HASPI scores with increasing hearing loss severity. The combined effect of increased severity and low SNR lead to low HASPI scores,

and this effect is expected based on the broader hearing aid literature and the inclusion of a hearing loss model within HASPI (Kates & Arehart, 2021).

Comparative Analysis Using a Non-Intrusive Metric

In evaluating the performance of signal processing approaches, it is informative to consider both intelligibility and sound quality. Sound quality can be predicted by the HASPI or by alternative metrics that use non-intrusive methods. Therefore, we estimated sound quality using the HASPI and pMOS, a non-intrusive model (Diehl et al., 2022). By analyzing the data using both HASPI and pMOS, and then evaluating the correlation between them, we aimed to assess the extent to which a non-intrusive model can approximate perceptual quality as measured by a validated intrusive metric. This approach is valuable because while HASPI requires a clean reference signal often unavailable in real-world applications, pMOS can be applied directly to degraded signals without the need for a reference. Demonstrating a strong correlation between HASPI and pMOS supports the use of non-intrusive models as practical, scalable alternatives for assessing speech quality, especially in large-scale or real-time systems where reference signals are impractical to obtain. This comparison also helps validate the pMOS model's reliability in approximating human-like perceptual judgments under realistic acoustic conditions. Figure 7 provides a scatterplot comparing the raw values obtained from each

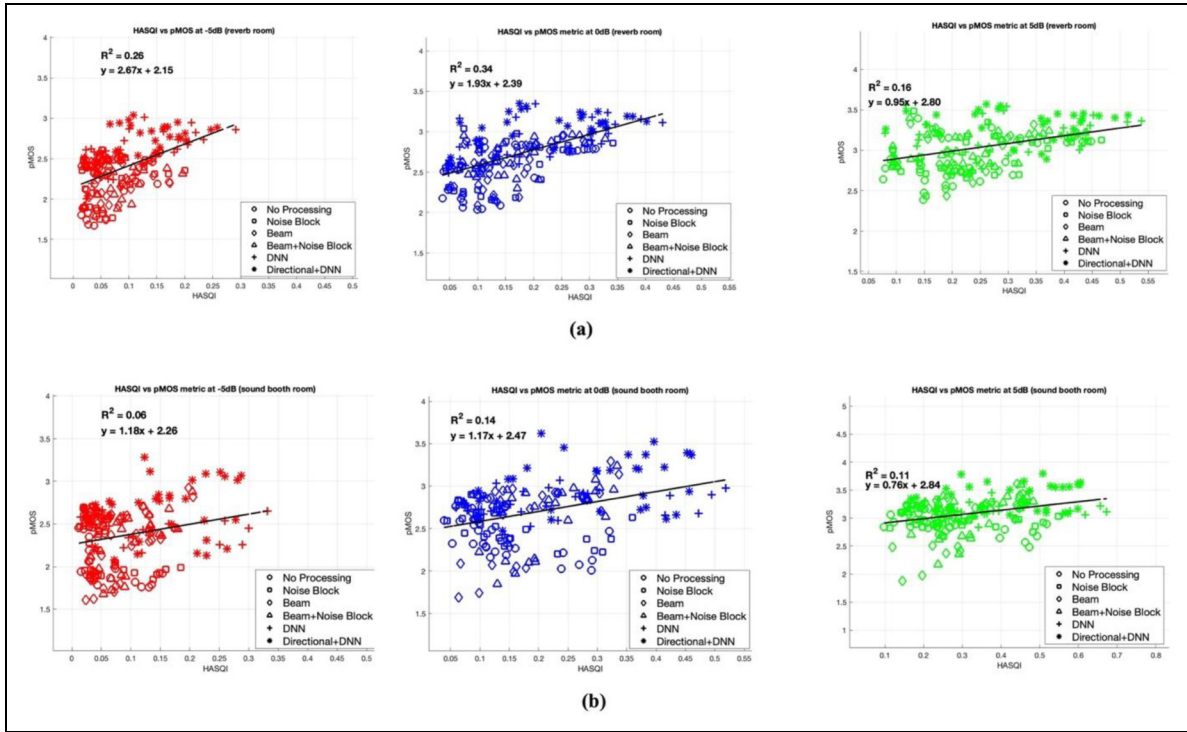


Figure 7. Raw Values for the HASQI and pMOS Non-Intrusive Metrics (Diehl et al., 2022), for Recordings Taken in the (a) Reverberant Room and (b) Sound Booth (Across all Conditions Mentioned in This Paper). Colors Indicate the Input SNR: Red: -5, Blue: 0, and Green: + 5 dB.

model, using the same recordings and signal processing settings described above for the HASPI analyses.

Figure 7 suggests a positive association between the HASQI scores and the pMOS metric across all conditions and processing methods. As the HASQI scores increase, indicating better speech quality, the pMOS metric also tends to increase, suggesting that both metrics generally agree on the relative performance of the processing methods. In both figures, in the -5 dB SNR condition (red symbols), both HASQI and pMOS scores are generally lower, but as the SNR increases to 0 dB (blue symbols) and 5 dB (green symbols), both HASQI and pMOS scores increase. This is consistent with sound quality estimates varying systematically with SNR, as would be expected. However, metrics from both the low and high reverberation environments indicate that there is some variability across processing conditions, where the pMOS metric shows a wider range of values. As one example, the beamforming plus DNN approach (asterisks) appears to yield higher pMOS scores compared to the HASQI scores. In contrast, the No Processing condition (circles) has a smaller range of variance. This relationship was further investigated by submitting these results to a hierarchical linear regression (SPSS v29), using two blocks and a step-wise method with variable entry and removal criteria of 5 and 10%, respectively. The HASQI values were entered as the predictor into the first block of a multiple linear regression, and with the pMOS values as the dependent (predicted)

variable. The correlation was significant ($F(1, 1150) = 1061.31$, $p < .001$), accounting for 47.9% of the variance. In the second block, the residual variance was the dependent variable, and the predictors were SNR, Azimuth, and Program Type in combination with the HASQI, accounting for 58.4% of the variance. These variables were added to the model in the order of SNR (adjusted R^2 change = 5.1%) and hearing aid program type (adjusted R^2 change = 5.3%), with each of these contributing significant changes in variance ($p < .001$). These results appear to indicate that the non-intrusive pMOS (Diehl et al., 2022) and HASQI metrics both index sound quality but respond differently to factors such as hearing aid program type and environmental variables.

Discussion

This study detailed the electroacoustic evaluation of a deep-neural-network-based denoising algorithm implemented in a hearing aid. Speech-in-noise recordings were collected with different hearing aid signal processing settings under a variety of fitting and environmental conditions. Speech intelligibility and quality predictors, namely HASPI and HASQI, were extracted from these recordings for benchmarking the signal processing feature performance. The HASQI and HASPI were chosen as they incorporate auditory models and are better suited for hearing aid assessments compared

to traditional metrics like PESQ and STOI (Kates & Arehart, 2014, 2021). Indeed, the HASPI and HASQI metrics were also the metrics of choice for performance evaluation in the recent clarity prediction and enhancement challenges for hearing aids (e.g., Akeroyd et al., 2023 and Barker et al., 2024). Additionally, a non-intrusive speech quality estimator, namely the pMOS, was also computed from the hearing aid recordings as it offers a promising alternative for evaluating speech quality in practical applications (Diehl et al., 2022).

Summary of DNN Effects

The findings of this study demonstrate potential advantages of DNN-based hearing aid processing, particularly when combined with beamforming, in improving speech intelligibility across diverse acoustic conditions (see Figures 2–6). Between 0- and 5-dB SNR, in a low reverberation environment, the DNN processor provided about 40–70% predicted speech intelligibility when combined with beamforming, while traditional noise reduction with and without beamforming provided only about 20–55% predicted speech intelligibility, across azimuths. This type of pattern was also observed in moderate reverberation, with better performance for more occluded fittings, for speech arising from the side, and better absolute performance for lesser degrees of hearing loss.

Some simulated cases had absolute predicted speech intelligibility above 80% in the beamforming plus DNN condition at a 5 dB SNR. Lower performance was associated with lower SNRs (below 20% at –5 dB SNR) and for more open fittings, with occluded fittings at about 50% versus about 20–35% for open and vented fittings. Consistent with previous findings (Andersen et al., 2021; Diehl et al., 2023) the DNN-based system enhanced predicted speech intelligibility for speech from a non-frontal location and at low SNRs. These improvements may address known limitations in traditional methods, including variable response times and reduced effectiveness at low SNRs, especially in reverberation (Ohlenforst et al., 2018; Reinhart et al., 2020; Scollie et al., 2016). Notably, the adaptability of the DNN algorithm to different noise types without relying on beamforming inputs suggests that it may offer an additional solution compared to beamforming, which in the past has primarily benefitted frontal sound sources (Kollmeier & Kiessling, 2018). Recent advances in adaptive and/or multiple parallel beamforming also show potential for improvement from multiple directions or multiple talkers (Andersson et al., 2021; Folkeard et al., 2024). Combination of DNN and multibeam adaptive beamforming may hold even further promise in the future.

The results demonstrate that DNN-based denoising consistently outperforms traditional adaptive filtering methods across different hearing aid fitting and listening conditions. It is pertinent to note that these benefits were observed even using the default DNN settings. Future research should investigate how adjusting the strength or configuration of DNN

features influences intelligibility and quality outcomes, as tuning may yield further improvements.

These findings contribute to the growing body of evidence supporting the use of deep learning for noise reduction in hearing aids (Healy et al., 2021; Goehring et al., 2016; Zhao et al., 2018). In particular, they suggest that DNN processing can improve performance without requiring complex spatial assumptions. Beyond denoising, other signal processing strategies such as acoustic scene classification and automatic program switching also play a key role in real-world noise management. Prior work has shown the value of such approaches (Jiang et al., 2020; Ye et al., 2018), and future work could explore how these methods might be integrated with DNN-based systems for adaptive, context-aware processing.

Effects of Reverberation and Venting

As mentioned earlier, across SNRs, the combination of DNN processing with beamforming consistently outperformed traditional beamforming and adaptive filtering methods. The most pronounced benefits were observed at 0 dB and +5 dB SNR in low-reverberation environments, where predicted speech perception improvements reached up to 33%, based on the HASPI data. In medium-reverberation conditions, although the DNN-based approach was less effective at the lowest SNR (–5 dB), it continued to yield significant benefits at higher SNRs. Analysis across venting conditions further demonstrated that the beamforming plus DNN configuration achieved high objective scores, with improved performance associated with reduced venting-consistent with trends observed in other hearing aid signal processing strategies.

However, the beamforming plus DNN conditions appear to provide higher metric values for predicted speech intelligibility across a broader range of vent sizes, particularly in the fully occluded condition. This advantage was observed even under 1 mm venting conditions. Previous research (Magnusson et al., 2013) indicates that the additive speech reception threshold benefits of combining beamforming with traditional adaptive noise reduction are only significant under occluded conditions. The current findings suggest that DNN-based denoising may extend benefit in noise to a broader range of fitting configurations, by allowing beamforming and denoising to work effectively together across a broader range of venting conditions.

Effects of Noise Type, Azimuth, and Audiogram

Further analyses across noise types, azimuths, and audiograms evaluated the impact of the DNN-based approach across listening situations and degrees of hearing loss. While performance was reduced in babble relative to performance in steady-state noise, the beamforming plus DNN still outperformed traditional methods by up to 30%. Beamforming benefits were evident, with the beamforming

plus DNN condition providing greater advantages at 0° and 90° azimuths, especially at higher SNRs. However, the magnitude of these benefits varied with hearing loss severity. Cases with more severe hearing loss (N4, N5 audiograms) demonstrated consistently lower HASPI scores across all conditions, with particularly poor performance at low SNRs, suggesting that the effectiveness of noise reduction and beamforming processing may be limited by auditory capacity. This result is expected based on the broader literature and because of the specific hearing loss model implemented in the HASPI (Kates & Arehart, 2021), which simulates reduced audibility, broadened auditory filters, and impaired temporal envelope processing. These factors likely contribute to diminished perceptual benefit from both noise reduction and beamforming processing in listeners with more severe hearing loss.

For speech from the side in a surround of noise, the DNN-alone system provided significantly better performance than traditional noise reduction and non-adaptive beamforming algorithms. This suggests that DNN denoising may enhance the ability to perceive speech from side angles without relying solely on beamforming microphone processing. This finding may be important for real-world listening environments where speech sources are not always positioned in front of the listener. Interestingly, the strategy with no signal processing performed better in the 90° condition than in the zero-degree condition. Recall that we selected the better-ear metric for all analyses. Therefore, the relative differences between these two azimuths likely represent better metric values in the ear receiving head shadow and may or may not reflect how a human listener performs in this condition.

While the current model-based results offer promising insights, further investigation using a larger number of test cases and behavioral data from human listeners is needed to determine whether individuals with greater degrees of hearing loss show increased candidacy for DNN-based processing. Importantly, our findings also revealed limitations in the predicted benefit of both traditional and DNN-based algorithms under challenging acoustic conditions—specifically in environments with moderate reverberation and in the presence of temporally modulated noise. These outcomes suggest that current DNN models may require further optimization to maintain effectiveness across a wider range of real-world listening environments. Nonetheless, the results underscore the potential of combining deep learning approaches with traditional signal processing strategies to improve speech perception, offering a foundation for the development of more robust and adaptable hearing aid technologies.

Sound Quality Predictions. This study also investigated the correlation between a non-intrusive speech quality metric, pMOS, and an intrusive speech quality metric, HASQI. While several non-intrusive estimators exist, the pMOS metric was chosen for comparison in this study because it was employed in training of the DNN architecture implemented within the tested hearing aids (Hasemann & Krylova, 2024).

It is pertinent to note here that the pMOS itself has been optimized to predict speech quality judgments from listeners with normal hearing (Diehl et al., 2022) and does not account for hearing loss effects. In contrast, HASQI integrates a comprehensive auditory peripheral model that simulates hearing loss and has been validated against subjective ratings from listeners with hearing impairment (Kates & Arehart, 2014; Kates & Arehart, 2022). We were therefore interested in understanding the relationship between these two metrics when the DNN operates alongside other hearing aid features, across a range of audiograms, hearing aid fitting configurations, and acoustic environments.

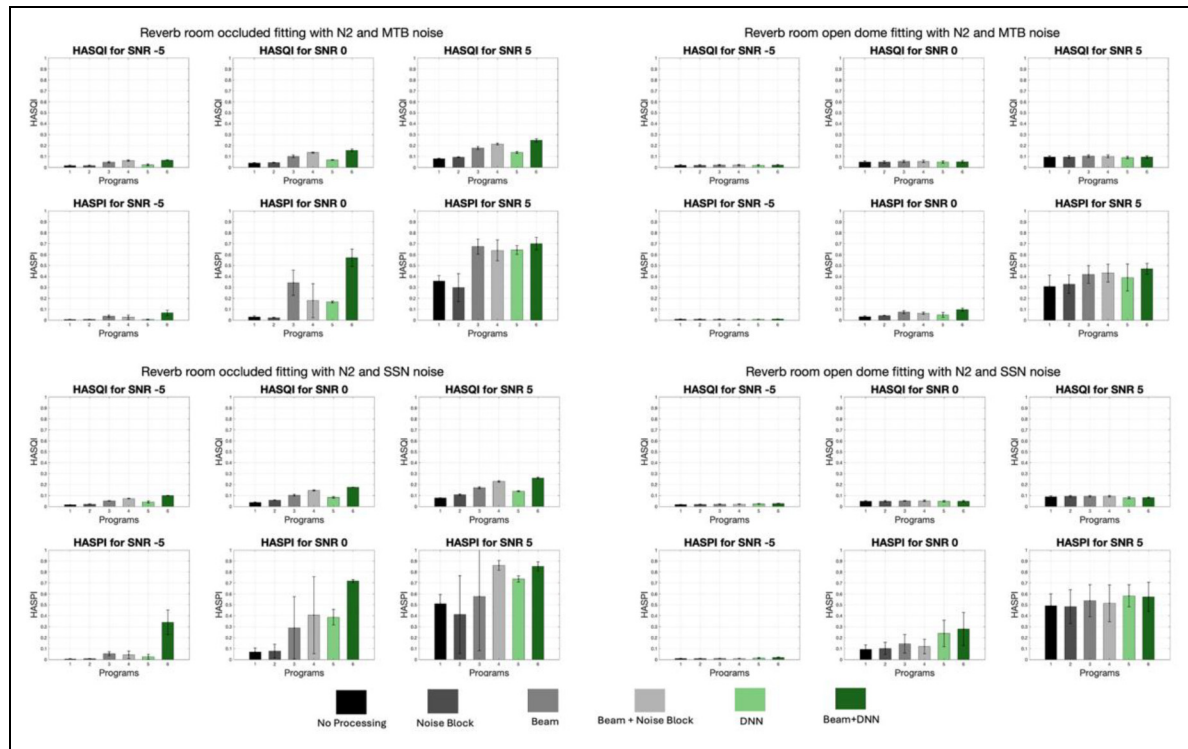
In general, the pMOS and HASQI metrics varied with one another, indicating cross-metric agreement across SNRs and across overall sound quality. They also varied in their response to specific hearing processing conditions and auditory scene factors. However, despite statistically significant correlations, the strength of the linear relationship between pMOS and HASQI varied considerably across different SNRs and reverberant environments, as reflected in the fluctuating amounts of variance explained by the regression models. These findings suggest that while both the non-intrusive pMOS metric and the intrusive HASQI metric capture aspects of sound quality, they exhibit distinct sensitivities to variables such as hearing aid program settings and environmental conditions. Further validation or extension of pMOS is needed to improve its ability to predict speech quality for hearing aid users.

Beyond pMOS, more recent research efforts have focused on developing improved non-intrusive metrics for hearing aid speech quality using speech foundation models (SFM), which leverage large-scale, pre-trained speech feature embeddings to support downstream tasks such as automatic speech recognition (Ahlawat et al., 2025). For instance, (Cuervo & Marxer, 2024) utilized SFMs to predict HASPI values from a database of hearing aid recordings, while (Chiang et al., 2024) developed HASA-Net+, an SFM-based model trained to estimate HASPI and HASQI values from simulated hearing aid data. More recently, (Ashkanichenarlogh et al., 2025) have developed a SFM based structure using automatic speech recognition and self-supervised learning trained using HASPI and HASQI values collected by commercialized HAs to predict speech quality and intelligibility indices for commercialized hearing aids, trained on custom noisy speech datasets and validated against intrusive measures and listener data. Future research could explore the use of these advanced non-intrusive metrics to evaluate hearing aids incorporating DNN denoising, ensuring that assessment methods keep pace with technological advancements.

Study Limitation. A key limitation of our findings lies in the reliance on objective metrics as proxies for subjective intelligibility and quality, particularly in the context of noise-reduction processing. While objective metrics such as HASPI, and HASQI are widely used due to their convenience

Table 3. Occluded and Vented Averaged Results in Different Acoustic Conditions Using Different Processing Methods.

	Acoustic environments	Speech from	No Processing (%)	Noise Block (%)	Beam (%)	Beam + Noise Block (%)	DNN (%)	Beam + DNN (%)	
HASQI	Reverb	0°	Occluded	9.36	10.6	21.05	23.25	14.41	26.36
			Vented	8.14	8.82	13.39	13.71	10.79	16.28
		90°	Occluded	21.41	23.69	11.83	14.13	30.86	29.11
			Vented	18.61	19.09	11.43	12.70	22.24	20.13
	Sound Booth	0°	Occluded	10.35	11.80	23.60	25.41	18.73	33.81
			Vented	11.38	12.04	18.57	18.85	15.79	24.15
		90°	Occluded	18.71	20.72	18.55	20.53	29.53	33.73
			Vented	18.74	18.83	17.68	18.82	22.77	25.76
HASPI	Reverb	0°	Occluded	11.16	10.55	29.67	30.86	25.35	44.07
			Vented	10.90	10.77	20.73	19.36	18.06	28.04
		90°	Occluded	30.39	34.47	14.74	16.85	51.06	45.97
			Vented	29.97	30.64	17.40	20.06	37.76	33.09
	Sound Booth	0°	Occluded	8.80	8.80	29.49	28.54	24.92	44.72
			Vented	9.80	10.54	23.71	22.64	18.84	31.47
		90°	Occluded	23.13	22.61	30.25	23.49	38.71	50.28
			Vented	21.25	21.73	27.37	27.78	29.37	37.91

**Figure 8.** HASQI and HASPI Scores in Reverberant Room for Speech in 0° for the N2 Audiogram with Different Conditions. Reverber room results for speech in 0°:

and reproducibility, recent evidence suggests that these metrics may not reliably predict actual human perception, especially when evaluating speech signals enhanced through deep learning methods. Notably, Gelderblom et al. (2023) demonstrated that none of the commonly used objective intelligibility metrics reliably predicted subjective intelligibility

improvements for speech signals processed by both single- and multi-channel deep complex convolution recurrent network (DCCRN)-based systems. Similarly, López-Espejo et al. (2023) showed that even when intelligibility and quality metrics are used as objective/loss functions in the training of the end-to-end neural networks, the resulting enhanced speech

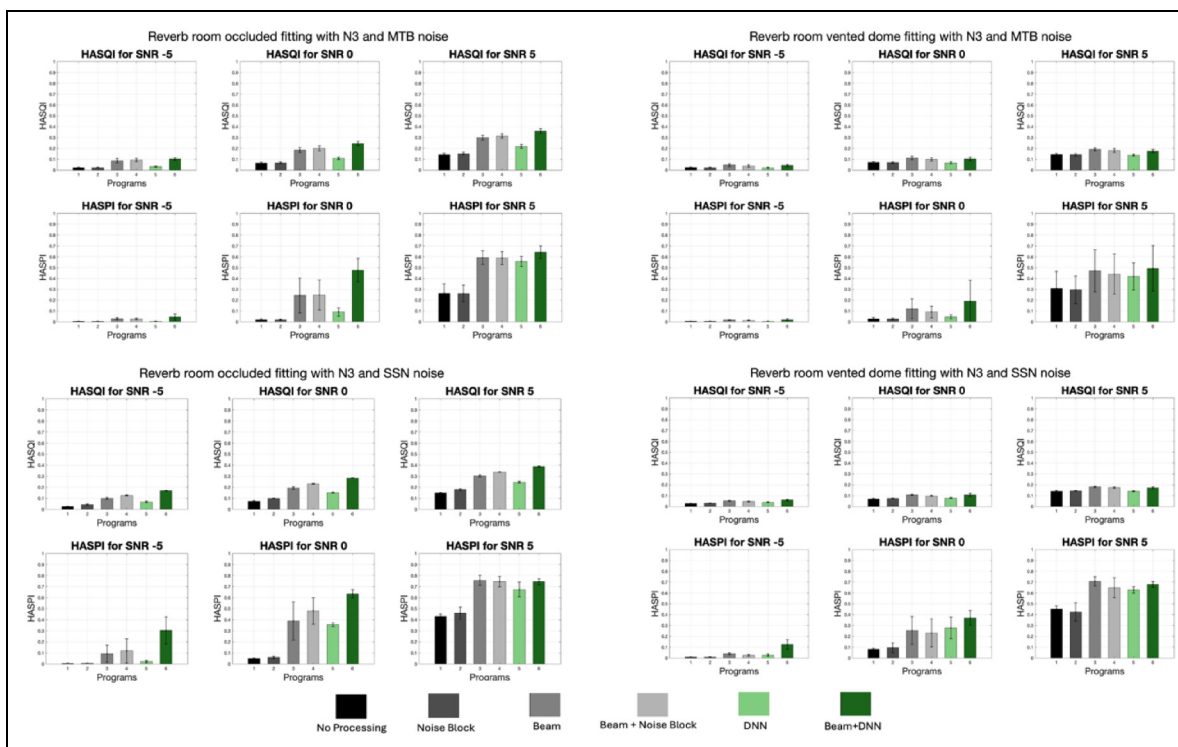


Figure 9. HASQI and HASPI Scores in Reverberant Room for Speech in 0° for the N3 Audiogram with Different Conditions.

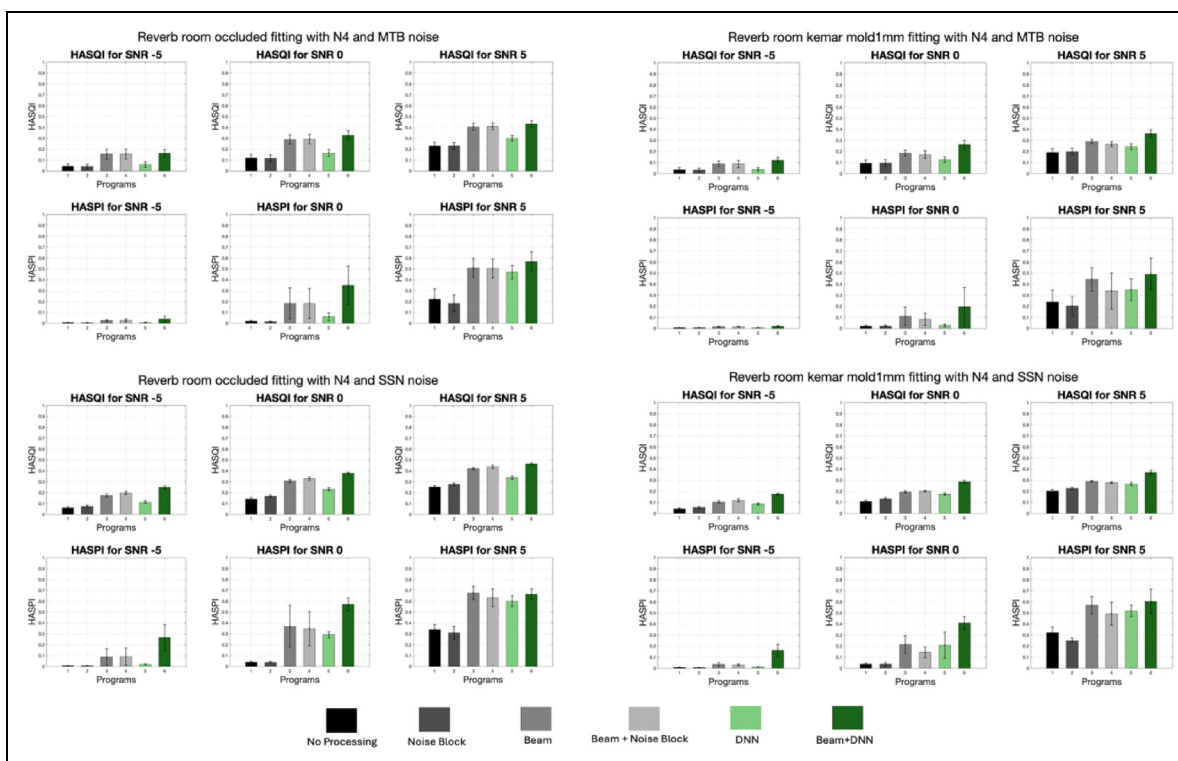


Figure 10. HASQI and HASPI Scores in Reverberant Room for Speech in 0° for the N4 Audiogram with Different Conditions.

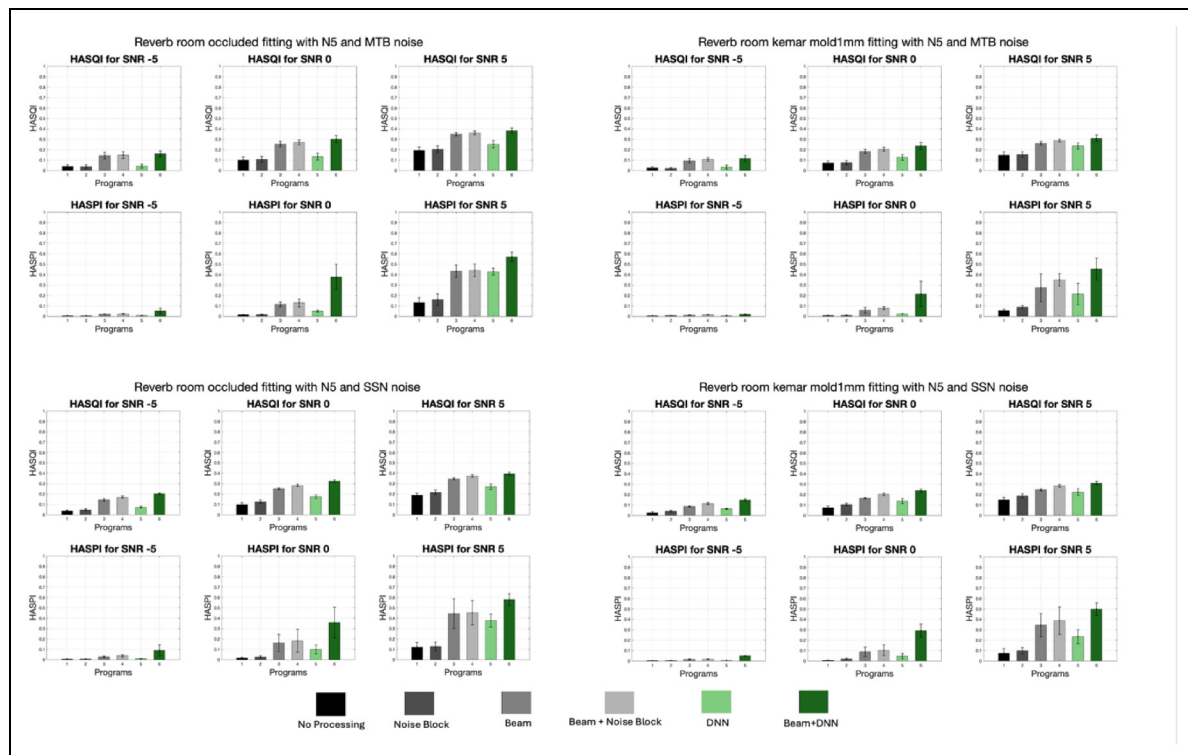


Figure 11. HASQI and HASPI Scores in Reverberant Room for Speech in 0° for the N5 Audiogram with Different Conditions. Reverber room results for speech in 90°:

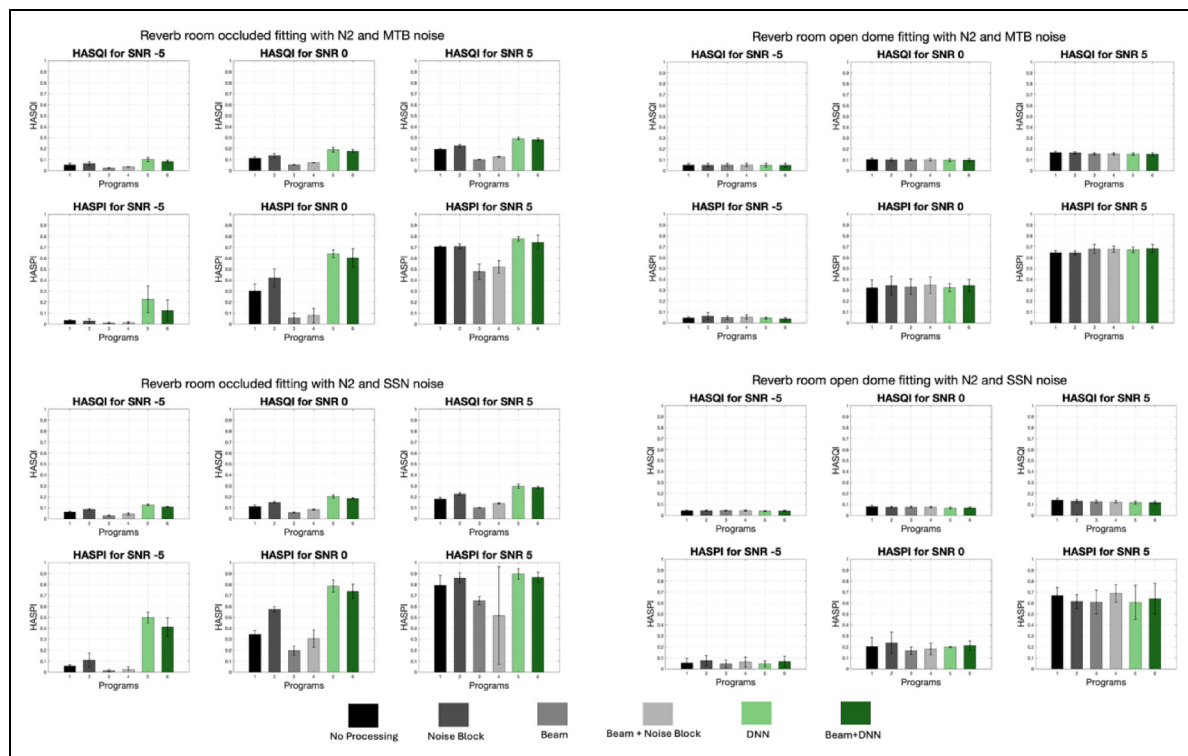


Figure 12. HASQI and HASPI Scores in Reverberant Room for Speech in 90° for the N2 Audiogram with Different Conditions.

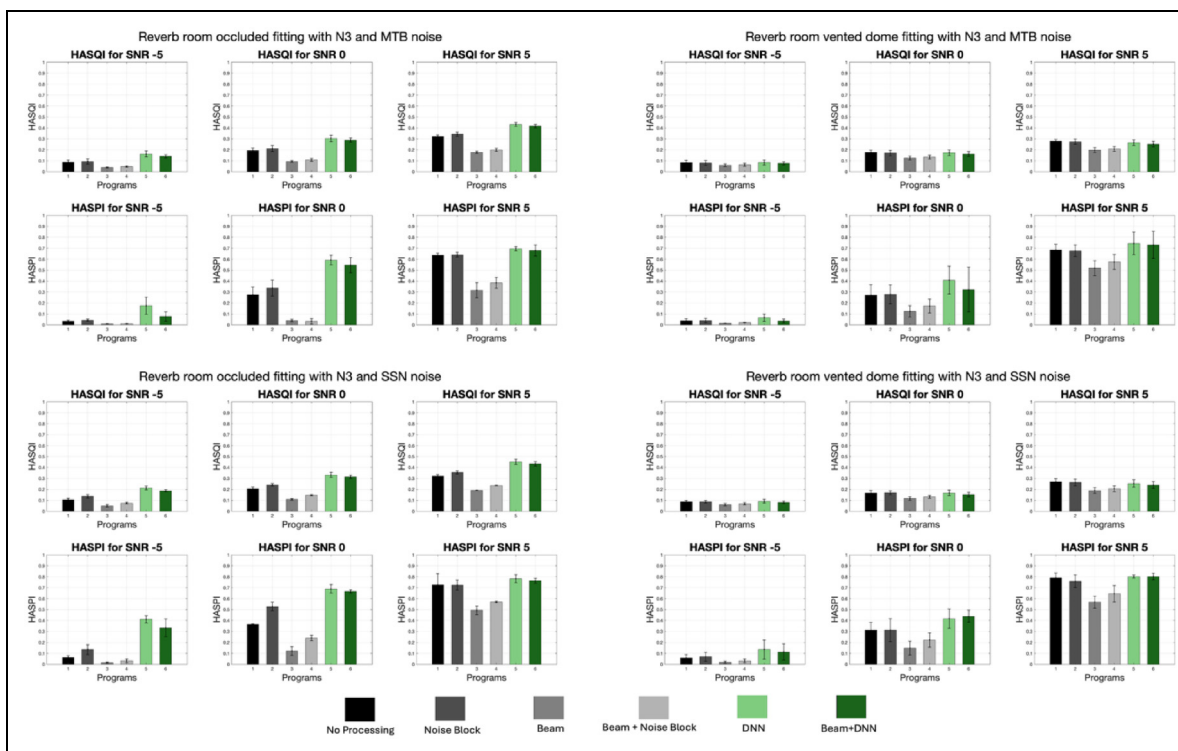


Figure 13. HASQI and HASPI Scores in Reverberant Room for Speech in 90° for the N3 Audiogram with Different Conditions.

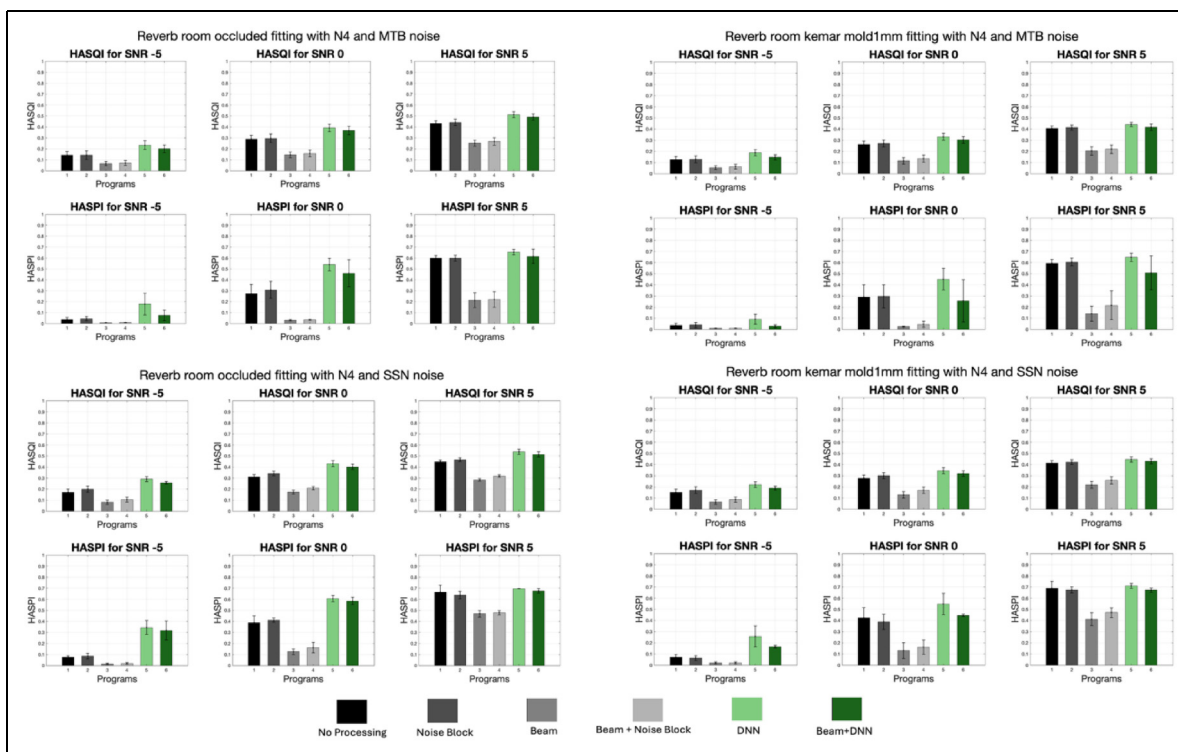


Figure 14. HASQI and HASPI Scores in Reverberant Room for Speech in 90° for the N4 Audiogram with Different Conditions.

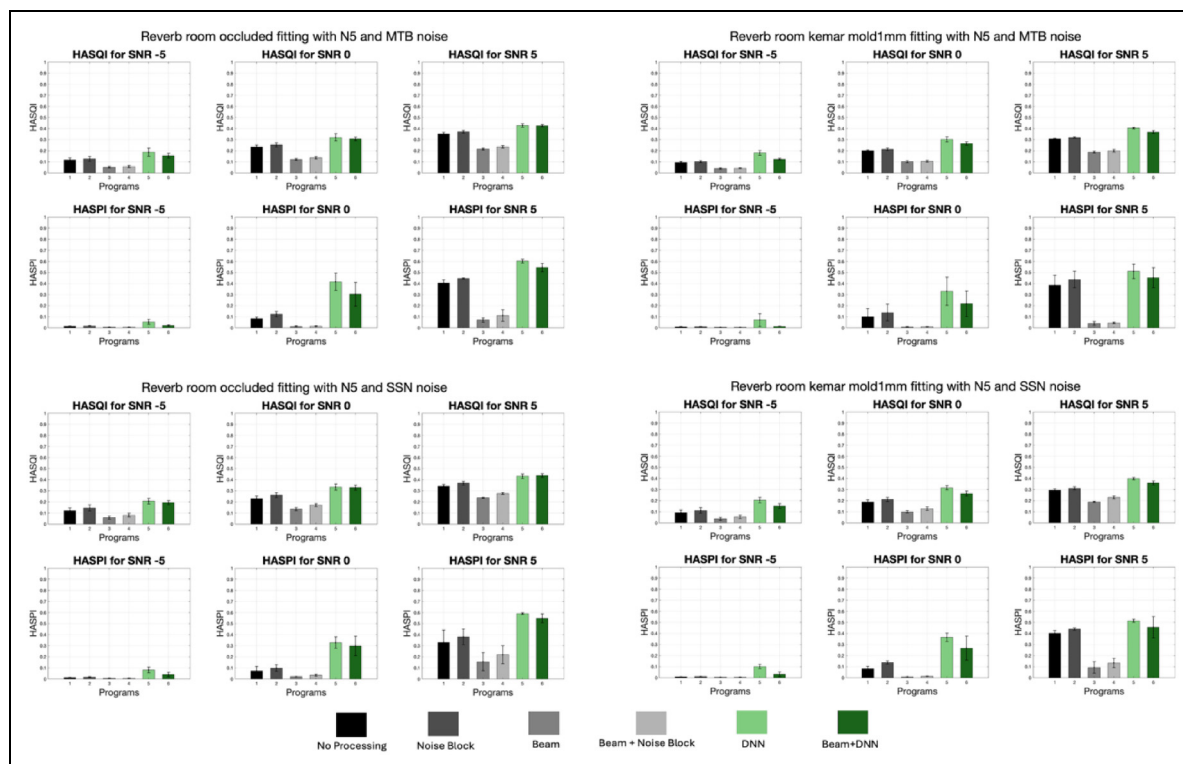


Figure 15. HASQI and HASPI Scores in Reverberant Room for Speech in 90° for the N5 Audiogram with Different Conditions. Sound booth results for speech in 0°:

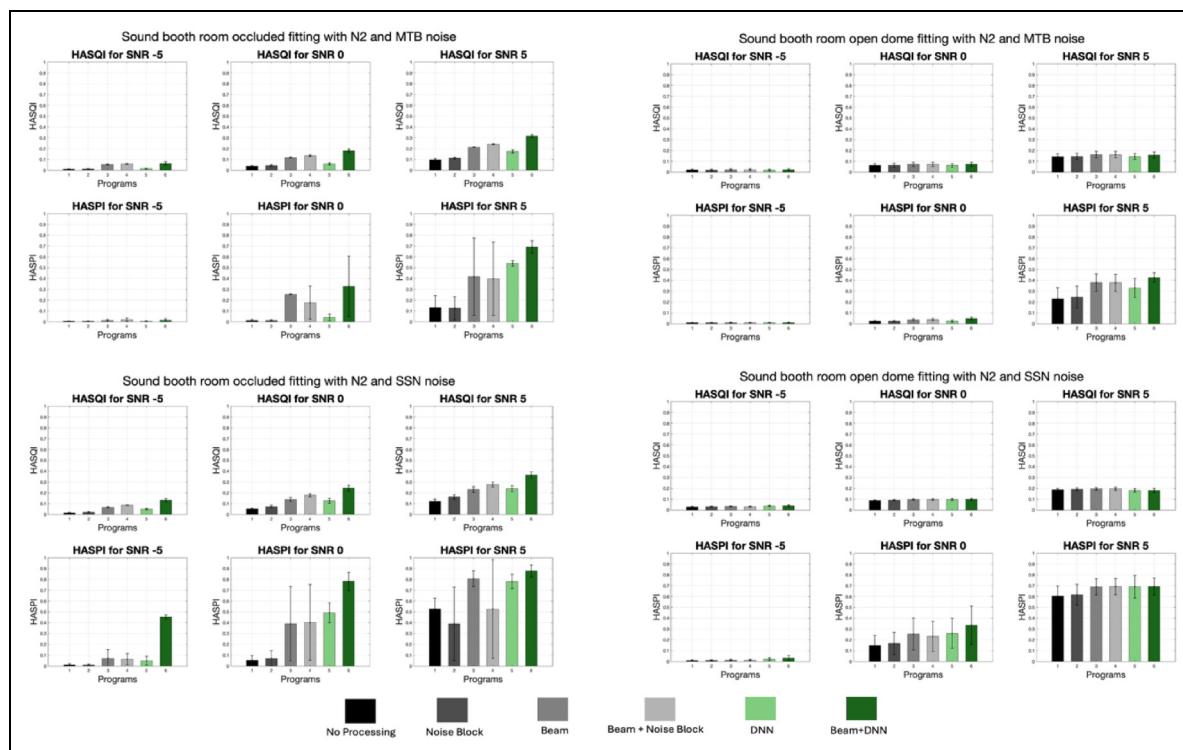


Figure 16. HASQI and HASPI Scores in Sound Booth Room for Speech in 0° for the N2 Audiogram with Different Conditions.

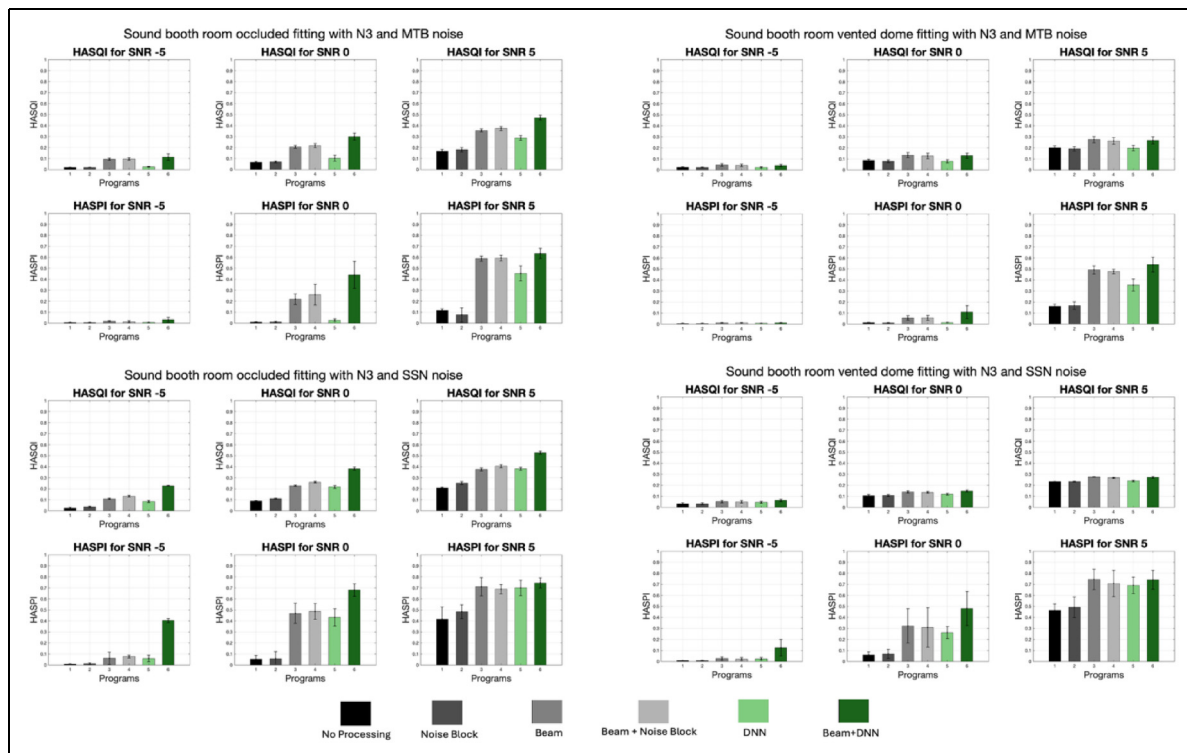


Figure 17. HASQI and HASPI Scores in Sound Booth Room for Speech in 0° for the N3 Audiogram with Different Conditions.

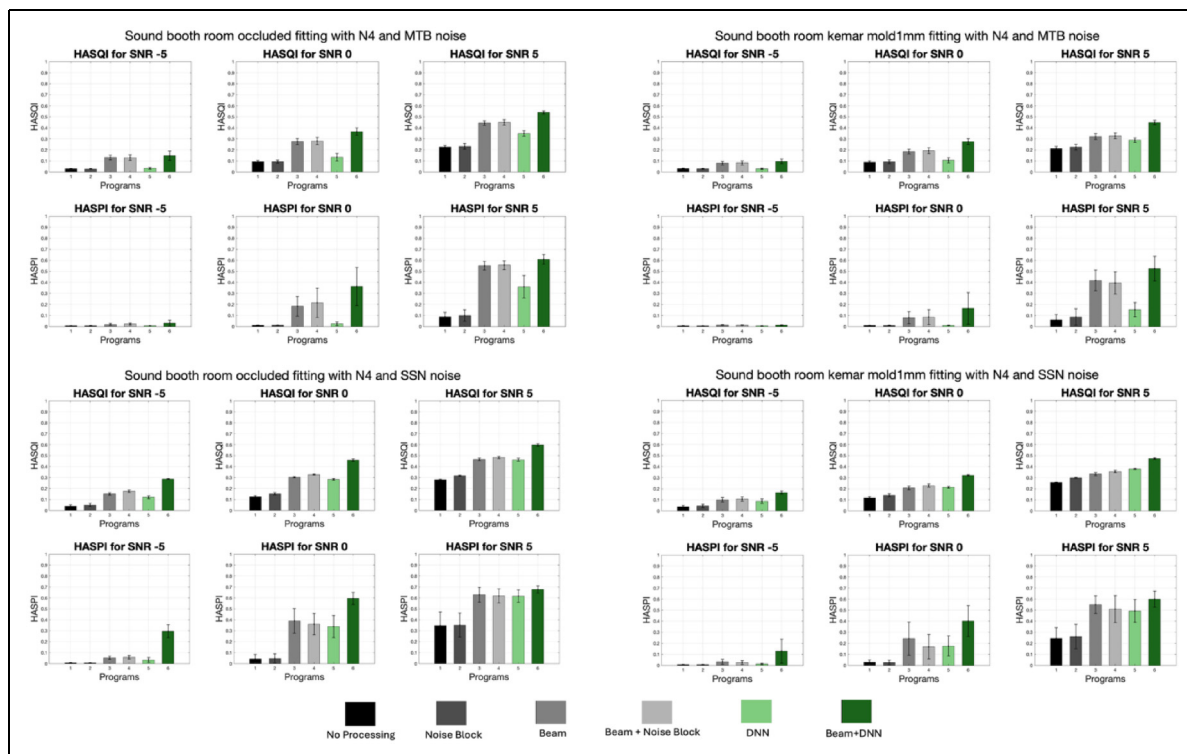


Figure 18. HASQI and HASPI Scores in Sound Booth Room for Speech in 0° for the N4 Audiogram with Different Conditions.

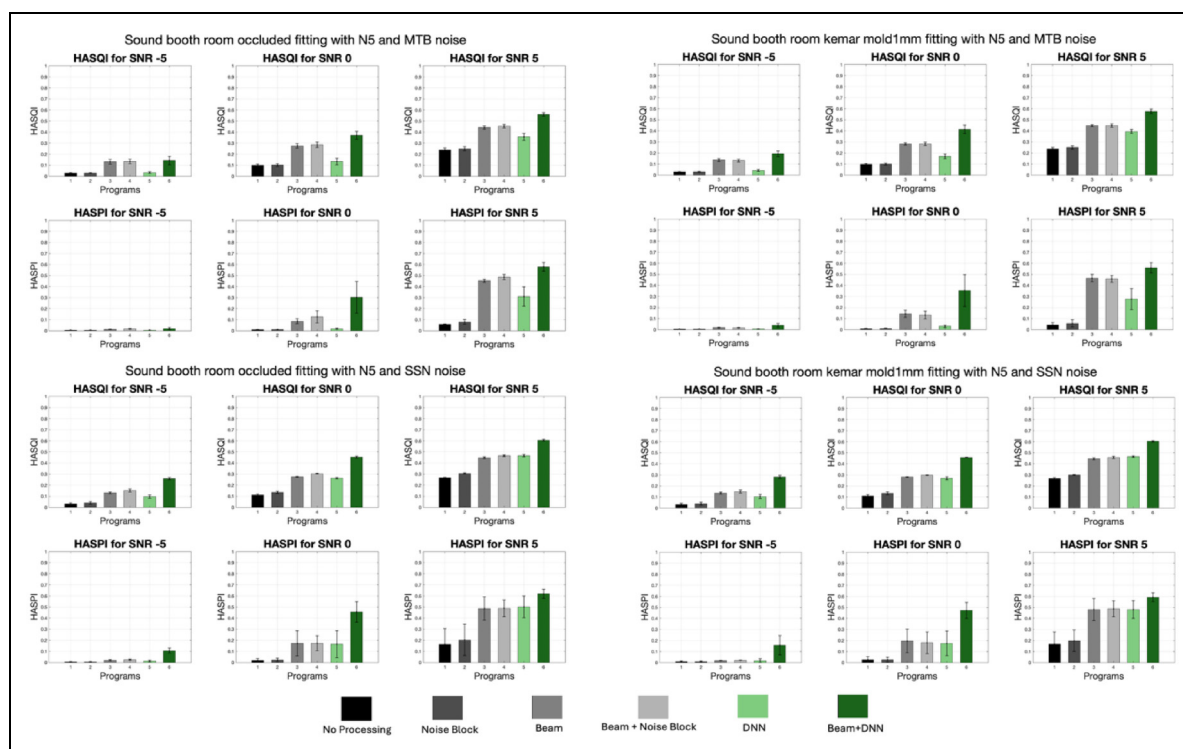


Figure 19. HASQI and HASPI Scores in Sound Booth Room for Speech in 0° for the N5 Audiogram with Different Conditions. Sound booth results for speech in 90° :

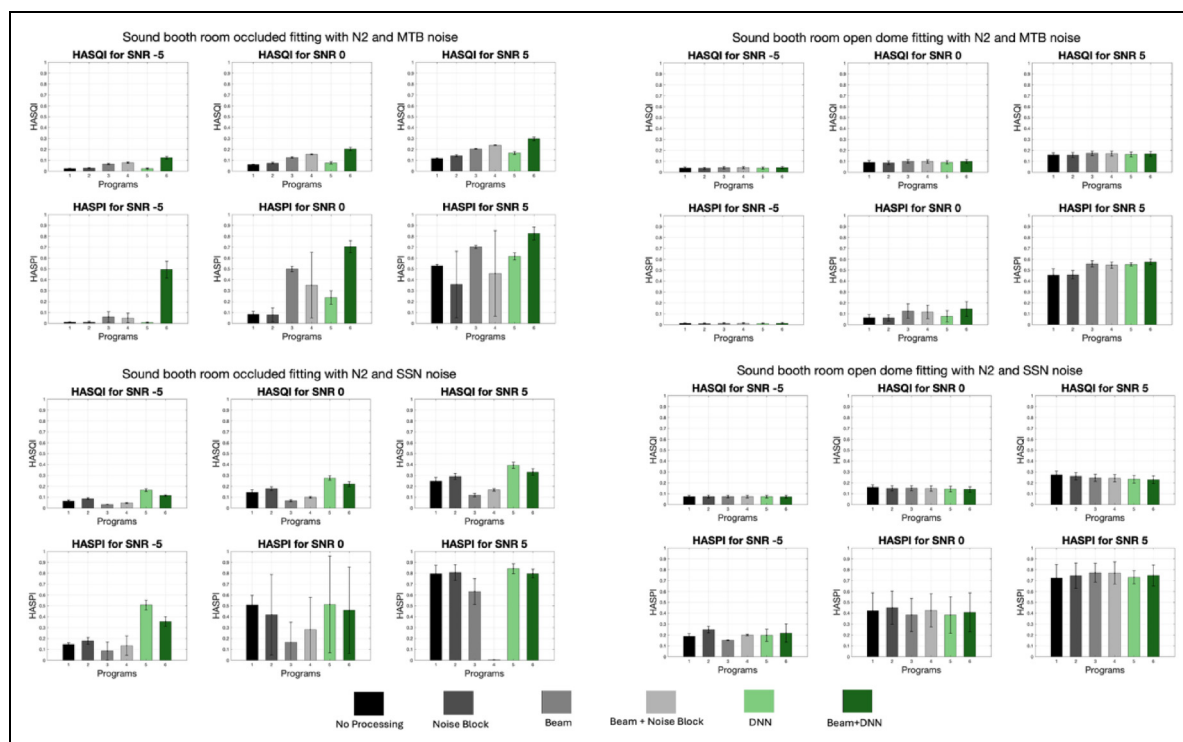


Figure 20. HASQI and HASPI Scores in Sound Booth Room for Speech in 90° for the N2 Audiogram with Different Conditions.

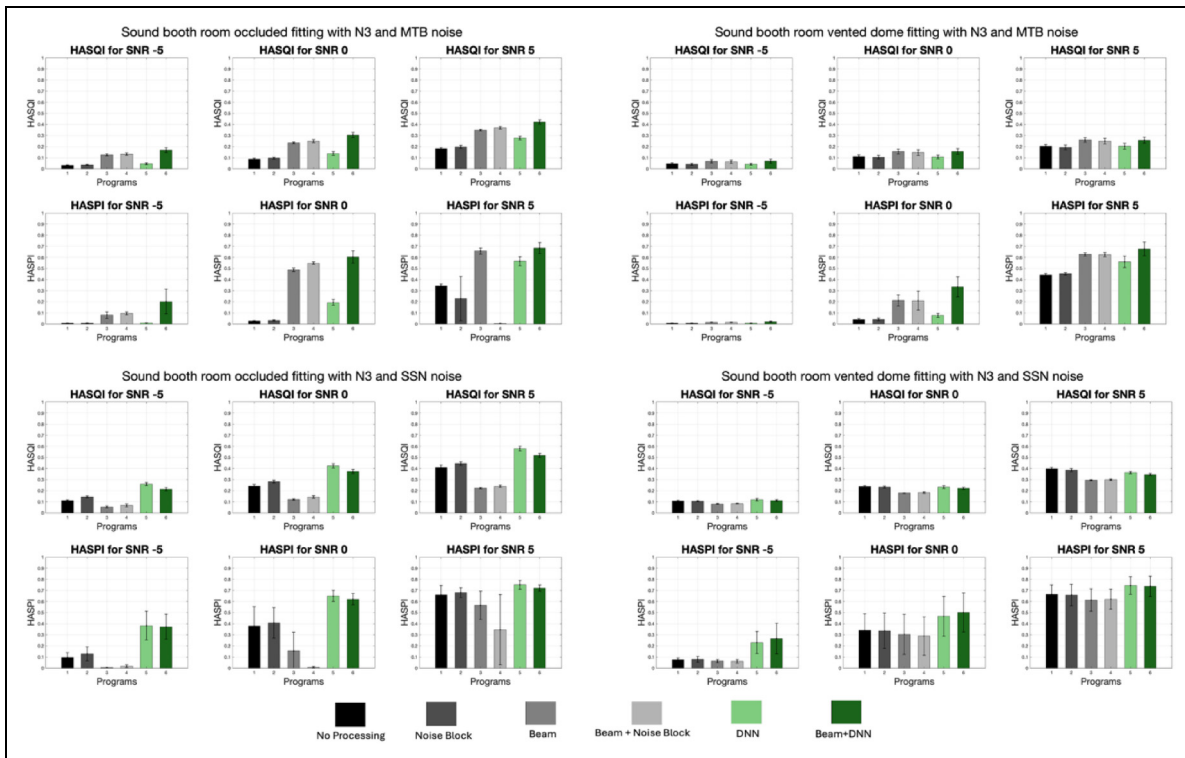


Figure 21. HASQI and HASPI Scores in Sound Booth Room for Speech in 90° for the N3 Audiogram with Different Conditions.

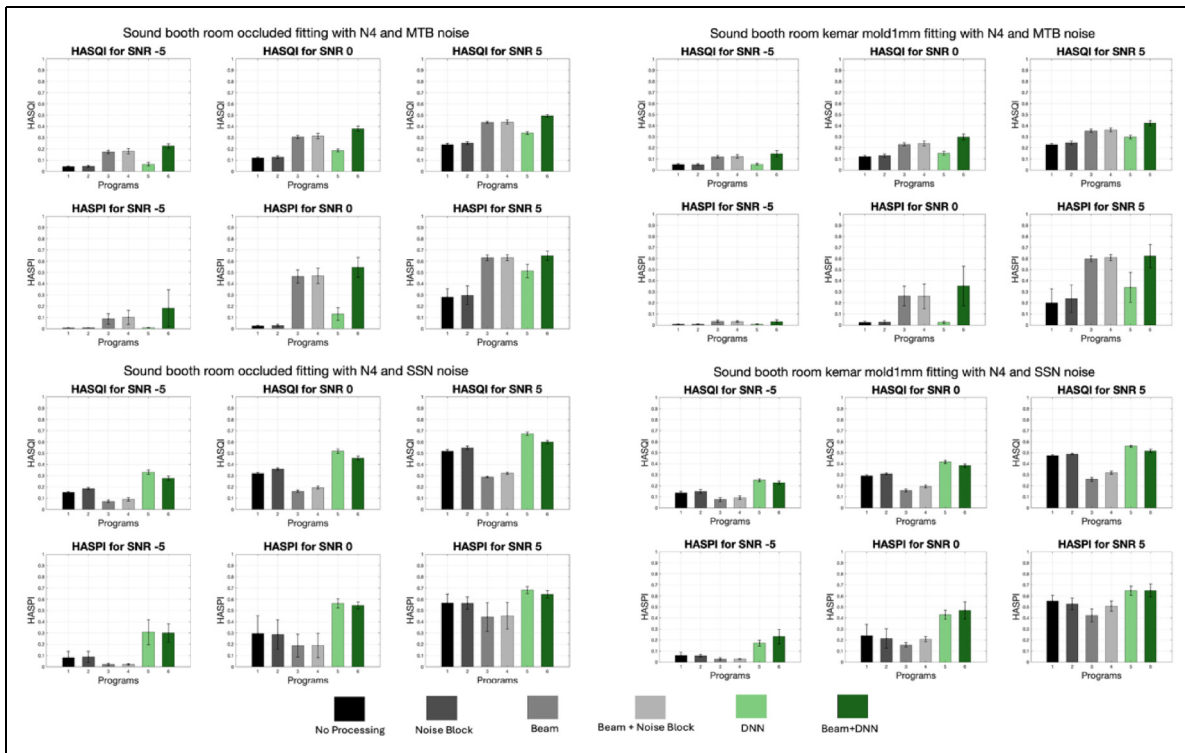


Figure 22. HASQI and HASPI Scores in Sound Booth Room for Speech in 90° for the N4 Audiogram with Different Conditions.

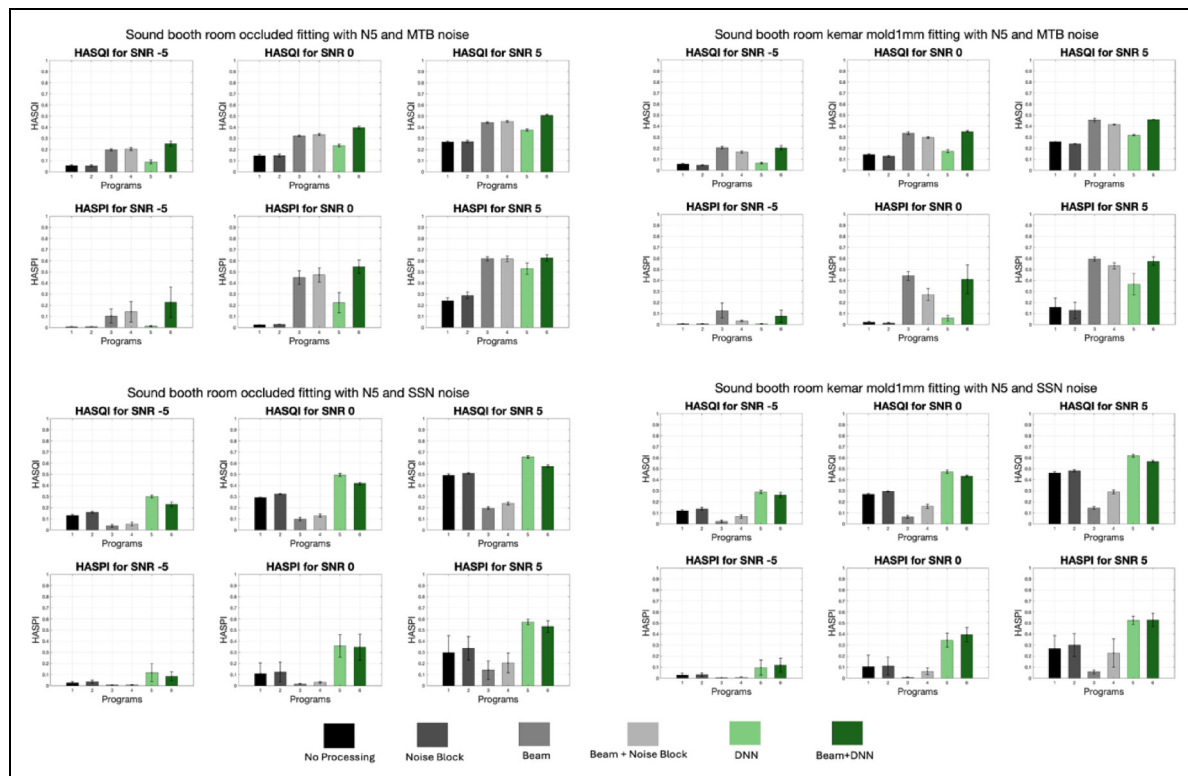


Figure 23. HASQI and HASPI Scores in Sound Booth Room for Speech in 90° for the N5 Audiogram with Different Conditions.

is not always characterized by better subjective intelligibility and quality. These findings suggest that such metrics may fail to fully capture perceptual dimensions introduced by modern processing approaches, including non-linear and context-sensitive transformations typical of DNNs. Thus, results suggest that while objective metrics remain useful for initial evaluations, their limitations must be acknowledged, and subjective listening tests should be incorporated when assessing real-world system performance. Consequently, while the present study provides a robust set of electroacoustic measurements and our findings show strong predicted benefits for DNN-based processing, these predictions must be interpreted cautiously and validated through behavioral testing with hard-of-hearing listeners. Such evaluations are essential for validating the real-world effectiveness of DNN denoising in improving speech intelligibility and sound quality. Gathering data from human listeners across hearing loss profiles will provide deeper insights into user experience and help refine DNN-based algorithms for broader clinical application.

Conclusion

This study demonstrates that implementing DNN-based denoising systems in commercially available hearing aids can substantially improve predicted speech intelligibility and sound quality across a wide range of listening conditions. Using objective auditory-model metrics (HASPI for

intelligibility, HASQI for quality), the DNN approach consistently outperformed traditional adaptive filtering and beamforming-based noise reduction, with the greatest benefits observed at lower SNRs, in steady-state noise, and for non-frontal speech sources. The integration of DNN processing with beamforming yielded the highest predicted intelligibility scores, particularly under low-reverberation conditions, with occluded fittings, and at SNRs of 0 and +5 dB. The DNN-alone configuration also demonstrated measurable improvements over non-DNN programs, notably for lateral speech presentations, indicating its potential to mitigate limitations of beamforming in multi-directional listening environments. Although performance gains were attenuated in medium-reverberation and multi-talker babble conditions, they remained evident, underscoring both the promise and current constraints of the tested system. It is important to note that HASPI estimates are based on a better-ear metric; thus, the observed advantages at 90° azimuth are likely attributable to head-shadow effects and may not reflect superior bilateral listener performance. Significant interactions between program type and factors such as noise type, venting configuration, and azimuth indicate that the magnitude of DNN benefit depends on both the acoustic environment and the physical fitting of the device. These results suggest that DNN-based denoising may serve as a valuable signal-processing option for challenging listening conditions that are not optimally addressed by conventional methods.

While the pMOS non-intrusive quality metric and the intrusive HASQI showed similar trends across SNRs, differences in their sensitivity to program type and environmental factors underscore the need for careful metric selection in future evaluations. Importantly, as this work relied solely on objective measures, further behavioral studies with hearing-impaired listeners are essential to validate whether these predicted benefits translate into perceptual gains in everyday listening.

As an appendix, the rest of this section has been provided to give more comprehensive insight to the readers. In this section Table 3, demonstrates the averaged occluded and vented results in percentage to conclude about the effects of occluded and vented fittings on the performance of different processing methods. This table presents the results of various signal processing techniques on speech quality (HASQI) and intelligibility (HASPI) across different acoustic environments, including reverb and sound booth conditions, with speech coming from 0° and 90° angles. The data demonstrates that the combination beamforming plus DNN consistently enhances both HASQI and HASPI scores across all conditions. Notably, occluded conditions, which initially show lower speech quality and intelligibility, benefit the most from this advanced processing, with significant improvements observed, particularly in challenging environments like low reverberation condition. Moreover, in this section Figures 8–23 presents a visualization of the HASQI and HASPI scores using different signal processing programs for all audiograms (N2, N3, N4, and N5) separately in different noisy (SSN, and MTB) conditions with different SNRs for the speech in 0° scenarios in both a reverberant room and sound booth. In these figures the x-axis shows the program or signal processing methods, while the y-axis illustrates the HASQI or HASPI values.

Acknowledgements

We would like to express our gratitude to Bilal Sheikh, Jinyu Qian, and Mathumitha Sureshkumar, and Yan Jiang for the invaluable collaboration throughout the research, and Stefan Raufer for review of the content of this manuscript prior to submission.

ORCID iDs

Vahid Ashkanichenarlogh  <https://orcid.org/0000-0003-2074-9863>

Paula Folkeard  <https://orcid.org/0000-0003-3109-6730>

Susan Scollie  <https://orcid.org/0000-0002-2207-5279>

Volker Kühnel  <https://orcid.org/0009-0004-9059-1592>

Ethical Considerations

There were no human participants, thus there was no ethics.

Consent to Participate

Not applicable.

Consent for Publication

Not applicable.

Funding

The authors are grateful for financial support for this research, authorship, and publication from Sonova and the Canada Mitacs Accelerate (IT29442) internship program.

SONOVA AG, Mitacs Accelerate, (grant number IT29442).

Declaration of Conflicting Interest

This research was funded by a grant from Sonova awarded to author S.S., and a Mitacs Accelerate Internship (IT29442) awarded to author V.A. Author V.K. is an employee of Sonova AG.

Data Availability

Data from this study is available at a GitHub link for the purposes of peer review: <https://github.com/vahidashkani/DNN-in-vented-and-occluded-hearing-aid>

If the manuscript is published, this data will also be made available with the publication, in accordance with the journal's objectives and formatting requirements.

References

- Abdel-Jaber, H., Devassy, D., Al Salam, A., Hidaytallah, L., & EL-Amir, M. (2022). A review of deep learning algorithms and their applications in healthcare. *Algorithms*, 15(2), 71. <https://doi.org/10.3390/a15020071>
- Ahlawat, H., Aggarwal, N., & Gupta, D. (2025, Jan 6). Automatic Speech Recognition: A survey of deep learning techniques and approaches. *International Journal of Cognitive Computing in Engineering*, 6(1), 201–237. <https://doi.org/10.1016/j.ijcce.2024.12.007>
- Akeroyd, M. A., Bailey, W., Barker, J., Cox, T. J., Culling, J. F., Graetzer, S., Naylor, G., Podwińska, Z., & Tu, Z. (2023). The 2nd clarity enhancement challenge for hearing aid speech intelligibility enhancement: Overview and outcomes. In *ICASSP 2023-2023 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP) 2023 Jun 4* (pp. 1–5). IEEE.
- Andersen, A. H., Santurette, S., Pedersen, M. S., Alickovic, E., Fiedler, L., Jensen, J., & Behrens, T. (2021). Creating clarity in noisy environments by using deep learning in hearing aids. *Seminars in Hearing*, 42(03), 260–281. Thieme Medical Publishers, Inc. <https://doi.org/10.1055/s-0041-1735134>
- Andersson, K. E., Andersen, L. S., Christensen, J. H., & Neher, T. (2021, Mar). Assessing real-life benefit from hearing-aid noise management: SSQ12 questionnaire versus ecological momentary assessment with acoustic data-logging. *American Journal of Audiology*, 10;30(1), 93–104. https://doi.org/10.1044/2020_AJA-20-00042
- Appleton, J., & König, G. (2014). Improvement in speech intelligibility and subjective benefit with binaural beamformer technology. *Hearing Review*, 21(10), 40–42.
- Ashkanichenarlogh, V., Folkeard, P., & Parsa, V. (2025, Apr 6). Towards Clinically Feasible Nonintrusive Quality and Intelligibility Indices for Hearing Aids. In *ICASSP 2025 IEEE*

- International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, pp. 1–5.
- Ashkanichenarlogh, V., & Parsa, V. (2024, Mar 1). Performance analysis of a dilated attention fast GAN for speech enhancement. *The Journal of the Acoustical Society of America*, 155(3_Supplement), A339–A340. <https://doi.org/10.1121/10.0027747>
- Barker, J., Akeroyd, M. A., Bailey, W., Cox, T. J., Culling, J. F., Firth, J., Graetzer, S., & Naylor, G. (2024). The 2nd Clarity Prediction Challenge: A machine learning challenge for hearing aid intelligibility prediction. In *ICASSP 2024-2024 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP) 2024 Apr 14* (pp. 11551–11555).
- Bentler, R. A., & Chiou, L. K. (2006). Digital noise reduction: An overview. *Trends in Amplification*, 10(2), 67–82. <https://doi.org/10.1177/1084713806289514>
- Bentsen, T., May, T., Kressner, A. A., & Dau, T. (2018). The benefit of combining a deep neural network architecture with ideal ratio mask estimation in computational speech segregation to improve speech intelligibility. *PLoS One*, 13(5), e0196924. <https://doi.org/10.1371/journal.pone.0196924>
- Bisgaard, N., Vlaming, M. S., & Dahlquist, M. (2010, Jun). Standard audiograms for the IEC 60118-15 measurement procedure. *Trends in Amplification*, 14(2), 113–120. <https://doi.org/10.1177/1084713810379609>
- 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>
- Borjigin, A., Kokkinakis, K., Bharadwaj, H. M., & Stohl, J. S. (2024). Deep learning restores speech intelligibility in multi-talker interference for cochlear implant users. *Scientific Reports*, 14(1), 13241. <https://doi.org/10.1038/s41598-024-63675-8>
- Boymans, M., & Dreschler, W. A. (2000, Jan). Field trials using a digital hearing aid with active noise reduction and dual-microphone directionality: Estudios de campo utilizando un audifono digital con reduccion activa del ruido y micrófono de direccionalidad dual. *Audiology*, 1;39(5), 260–268. <https://doi.org/10.3109/002060900009073090>
- Bramsløw, L., Naithani, G., Hafez, A., Barker, T., Pontoppidan, N. H., & Virtanen, T. (2018, Jul 1). Improving competing voices segregation for hearing impaired listeners using a low-latency deep neural network algorithm. *The Journal of the Acoustical Society of America*, 144(1), 172–185. <https://doi.org/10.1121/1.5045322>
- Brons, I., Houben, R., & Dreschler, W. A. (2014, Oct 13). Effects of noise reduction on speech intelligibility, perceived listening effort, and personal preference in hearing-impaired listeners. *Trends in Hearing*, 18, 2331216514553924. <https://doi.org/10.1177/2331216514553924>
- Brons, I., Houben, R., & Dreschler, W. A. (2014). Perceptual effects of noise reduction with respect to personal preference, speech intelligibility, and listening effort. *Ear and Hearing*, 35(5), 519–530. <https://doi.org/10.1097/AUD.0000000000000021>
- Chiang, H. T., Fu, S. W., Wang, H. M., Tsao, Y., & Hansen, J. H. (2024, Nov 1). Multi-objective non-intrusive hearing-aid speech assessment model. *The Journal of the Acoustical Society of America*, 156(5), 3574–3587. <https://doi.org/10.1121/10.0034362>
- Chong, F. Y., & Jenstad, L. M. (2018, Aug 18). A critical review of hearing-aid single-microphone noise-reduction studies in adults and children. *Disability and Rehabilitation: Assistive Technology*, 13(6), 600–608. <https://doi.org/10.1080/17483107.2017.1392619>
- Christensen, J. H., Whiston, H., Lough, M., Gil-Carvajal, J. C., Rumley, J., & Saunders, G. H. (2024, Mar 4). Evaluating real-world benefits of hearing aids with deep neural network-based noise reduction: An ecological momentary assessment study. *American Journal of Audiology*, 3(1), 242–253. https://doi.org/10.1044/2023_AJA-23-00149
- Collins, G. S., Reitsma, J. B., Altman, D. G., & Moons, K. G. (2015, Jan 13). Transparent reporting of a multivariable prediction model for individual prognosis or diagnosis (TRIPOD) the TRIPOD statement. *Circulation*, 131(2), 211–219. <https://doi.org/10.1161/CIRCULATIONAHA.114.014508>
- Cuervo, S., & Marxer, R. (2024, Apr 14). Speech foundation models on intelligibility prediction for hearing-impaired listeners. In *ICASSP 2024-2024 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, pp. 1421–1425.
- Desjardins, J. L. (2016 Jan). The effects of hearing aid directional microphone and noise reduction processing on listening effort in older adults with hearing loss. *Journal of the American Academy of Audiology*, 27(01), 029–041. <https://doi.org/10.3766/jaaa.15030>
- Diehl, P. U., Singer, Y., Zilly, H., Schönfeld, U., Meyer-Rachner, P., Berry, M., Sprekeler, H., Sprengel, E., Pudzuhn, A., & Hofmann, V. M. (2023 Feb 15). Restoring speech intelligibility for hearing aid users with deep learning. *Scientific Reports*, 13(1), 2719. <https://doi.org/10.1038/s41598-023-29871-8>
- Diehl, P. U., Thorbergsson, L., Singer, Y., Skripniuk, V., Pudzuhn, A., Hofmann, V. M., Sprengel, E., & Meyer-Rachner, P. (2022, Nov 28). Non-intrusive deep learning-based computational speech metrics with high-accuracy across a wide range of acoustic scenes. *Plos one*, 17(11), e0278170. <https://doi.org/10.1371/journal.pone.0278170>
- Falk, T. H., Parsa, V., Santos, J. F., Arehart, K., Hazrati, O., Huber, R., Kates, J. M., & Scollie, S. (2015, Feb 10). Objective quality and intelligibility prediction for users of assistive listening devices: Advantages and limitations of existing tools. *IEEE signal Processing Magazine*, 32(2), 114–124. <https://doi.org/10.1109/MSP.2014.2358871>
- Folkeard, P., Jensen, N. S., Parsi, H. K., Bilert, S., & Scollie, S. (2024, Sep 3). Hearing at the mall: Multibeam processing technology improves hearing group conversations in a real-world food court environment. *American Journal of Audiology*, 33(3), 782–792. https://doi.org/10.1044/2024_AJA-24-00027
- Gaultier, C., & Goehring, T. (2024). Recovering speech intelligibility with deep learning and multiple microphones in noisy-reverberant situations for people using cochlear implants. *The Journal of the Acoustical Society of America*, 155(6), 3833–3847. <https://doi.org/10.1121/10.0026218>
- Gelderblom, F. B., Tronstad, T. V., Svendsen, T., & Myrvoll, T. A. (2023). On the predictive power of objective intelligibility metrics for the subjective performance of deep complex

- convolutional recurrent speech enhancement networks. *IEEE/ACM Trans. Audio, Speech, Lang. Process.*, 32, 215–226. <https://doi.org/10.1109/TASLP.2023.3329378>
- Goehring, T., Bolner, F., Monaghan, J. J., Van Dijk, B., Zarowski, A., & Bleeck, S. (2017, Feb 1). Speech enhancement based on neural networks improves speech intelligibility in noise for cochlear implant users. *Hearing Research*, 344, 183–194. <https://doi.org/10.1016/j.heares.2016.11.012>
- Goehring, T., Yang, X., Monaghan, J. J., & Bleeck, S. (2016, Aug 29). Speech enhancement for hearing-impaired listeners using deep neural networks with auditory model-based features. In *2016 24th European signal processing conference (EUSIPCO)*, pp. 2300–2304.
- Hasemann, H., & Krylova, A. (April 2024). Revolutionary Speech Understanding with Spheric Speech Clarity, Phonak Insight, pp. 1–7.
- Healy, E. W., Johnson, E. M., Delfarah, M., Krishnagiri, D. S., Sevich, V. A., Taherian, H., & Wang, D. (2021, Oct 1). Deep learning-based speaker separation and dereverberation can generalize across different languages to improve intelligibility. *The Journal of the Acoustical Society of America*, 150(4), 2526–2538. <https://doi.org/10.1121/10.0006565>
- Healy, E. W., Johnson, E. M., Pandey, A., & Wang, D. (2023). Progress made in the efficacy and viability of deep-learning-based noise reduction. *The Journal of the Acoustical Society of America*, 153(5), 2751–2768. <https://doi.org/10.1121/10.0019341>
- Healy, E. W., Taherian, H., Johnson, E. M., & Wang, D. (2021). A causal and talker-independent speaker separation/dereverberation deep learning algorithm: Cost associated with conversion to real-time capable operation. *The Journal of the Acoustical Society of America*, 150(5), 3976–3986. <https://doi.org/10.1121/10.0007134>
- Healy, E. W., Tan, K., Johnson, E. M., & Wang, D. (2021, Jun 1). An effectively causal deep learning algorithm to increase intelligibility in untrained noises for hearing-impaired listeners. *The Journal of the Acoustical Society of America*, 149(6), 3943–3953. <https://doi.org/10.1121/10.0005089>
- Hornsby, B. W., & Ricketts, T. A. (2007). Effects of noise source configuration on directional benefit using symmetric and asymmetric directional hearing aid fittings. *Ear and Hearing*, 28(2), 177–186. <https://doi.org/10.1097/AUD.0b013e3180312639>
- Jiang, D., Huang, D., Song, Y., Wu, K., Lu, H., Liu, Q., & Zhou, T. (2020, Sep 28). An audio data representation for traffic acoustic scene recognition. *IEEE access*, 8, 177863–73. <https://doi.org/10.1109/ACCESS.2020.3027474>
- Kates, J. M., & Arehart, K. H. (2014, Mar 20). The hearing-aid speech quality index (HASQI) version 2. *Journal of the Audio Engineering Society*, 62(3), 99–117. <https://doi.org/10.17743/jaes.2014.0006>
- Kates, J. M., & Arehart, K. H. (2021, Jul 1). The hearing-aid speech perception index (HASPI) version 2. *Speech Communication*, 131, 35–46. <https://doi.org/10.1016/j.specom.2020.05.001>
- Kates, J. M., & Arehart, K. H. (2022). An overview of the HASPI and HASQI metrics for predicting speech intelligibility and speech quality for normal hearing, hearing loss, and hearing aids. *Hearing Research*, 426, 108608. <https://doi.org/10.1016/j.heares.2022.108608>
- Keidser, G., Carter, L., Chalupper, J., & Dillon, H. (2007, Jan 1). Effect of low-frequency gain and venting effects on the benefit derived from directionality and noise reduction in hearing aids. *International Journal of Audiology*, 46(10), 554–568. <https://doi.org/10.1080/14992020701481698>
- Keshavarzi, M., Goehring, T., Turner, R. E., & Moore, B. C. (2019). Comparison of effects on subjective intelligibility and quality of speech in babble for two algorithms: A deep recurrent neural network and spectral subtraction. *The Journal of the Acoustical Society of America*, 145(3), 1493–1503. <https://doi.org/10.1121/1.5094765>
- Kollmeier, B., & Kiessling, J. (2018, May 25). Functionality of hearing aids: State-of-the-art and future model-based solutions. *International Journal of Audiology*, 57(sup3), S3–28. <https://doi.org/10.1080/14992027.2016.1256504>
- Krishnamurthy, N., & Hansen, J. H. (2009, Jul 28). Babble noise: Modeling, analysis, and applications. *IEEE transactions on Audio, Speech, and Language Processing*, 17(7), 1394–1407. <https://doi.org/10.1109/TASL.2009.2015084>
- Krueger, M., Schulte, M., Zokoll, M. A., Wagener, K. C., Meis, M., Brand, T., & Holube, I. (2017, Oct 12). Relation between listening effort and speech intelligibility in noise. *American Journal of Audiology*, 26(3S), 378–392. https://doi.org/10.1044/2017_AJA-16-0136
- Kumar, S., Guruvayurappan, A., Pitchaimuthu, A. N., & Nayak, S. (2023, Nov 1). Efficacy and effectiveness of wireless binaural beamforming technology of hearing aids in improving speech perception in noise: A systematic review. *Ear and Hearing*, 44(6), 1289–1300. <https://doi.org/10.1097/AUD.0000000000001374>
- Lakshmi, M. S., Rout, A., & O'Donoghue, C. R. (2021, Feb 17). A systematic review and meta-analysis of digital noise reduction hearing aids in adults. *Disability and Rehabilitation: Assistive Technology*, 16(2), 120–129. <https://doi.org/10.1080/17483107.2019.1642394>
- Lansbergen, S., & Dreschler, W. A. (2022, Feb 21). Classification of hearing aids into feature profiles using hierarchical latent class analysis applied to a large dataset of hearing aids. *Ear and hearing*.
- Launer, S., Zakis, J. A., & Moore, B. C. (2016). Hearing aid signal processing. *Hearing Aids*, 56, 93–130. https://doi.org/10.1007/978-3-319-33036-5_4
- Lelic, D., Wolters, F., Herrlin, P., & Smeds, K. (2022, Dec 5). Assessment of hearing-related lifestyle based on the common sound scenarios framework. *American Journal of Audiology*, 31(4), 1299–1311. https://doi.org/10.1044/2022_AJA-22-00079
- López-Espejo, I., Edraki, A., Chan, W.-Y., Tan, Z.-H., & Jensen, J. (2023). On the deficiency of intelligibility metrics as proxies for subjective intelligibility. *Speech Communication*, 150, 9–22. <https://doi.org/10.1016/j.specom.2023.04.001>
- Magnusson, L., Claesson, A., Persson, M., & Tengstrand, T. (2013). Speech recognition in noise using bilateral open-fit hearing aids: The limited benefit of directional microphones and noise reduction. *International Journal of Audiology*, 52(1), 29–36. <https://doi.org/10.3109/14992027.2012.707335>
- Ohlenforst, B., Wendt, D., Kramer, S. E., Naylor, G., Zekveld, A. A., & Lunner, T. (2018, Aug 1). Impact of SNR, masker type and

- noise reduction processing on sentence recognition performance and listening effort as indicated by the pupil dilation response. *Hearing Research*, 365, 90–99. <https://doi.org/10.1016/j.heares.2018.05.003>
- Pichora-Fuller, M. K., Kramer, S. E., Eckert, M. A., Edwards, B., Hornsby, B. W., Humes, L. E., Lemke, U., Lunner, T., Matthen, M., Mackersie, C. L., & Naylor, G. (2016, Jul 1). Hearing impairment and cognitive energy: The framework for understanding effortful listening (FUEL). *Ear and Hearing*, 37, 5S–27S. <https://doi.org/10.1097/AUD.0000000000000312>
- Picou, E. M. (2022 Nov). Hearing aid benefit and satisfaction results from the MarkeTrak 2022 survey: Importance of features and hearing care professionals. *Seminars in Hearing*, 43(04), 301–316. Thieme Medical Publishers, Inc. <https://doi.org/10.1055/s-0042-1758375>
- Picou, E. M., Aspell, E., & Ricketts, T. A. (2014, May 1). Potential benefits and limitations of three types of directional processing in hearing aids. *Ear and Hearing*, 35(3), 339–352. <https://doi.org/10.1097/AUD.0000000000000004>
- Reinhart, P., Zahorik, P., & Souza, P. (2020, Jan). Interactions between digital noise reduction and reverberation: Acoustic and behavioral effects. *Journal of the American Academy of Audiology*, 31(01), 017–029. <https://doi.org/10.3766/jaaa.18048>
- Renz, T., Leistner, P., & Liebl, A. (2018, Mar 1). Auditory distraction by speech: Sound masking with speech-shaped stationary noise outperforms–5dB per octave shaped noise. *The Journal of the Acoustical Society of America*, 143(3), EL212–7. <https://doi.org/10.1121/1.5027765>
- Rhebergen, K. S., Versfeld, N. J., & Dreschler, W. A. (2010). Release from informational masking by time reversal of noise in speech recognition. *Journal of the Acoustical Society of America*, 127(1), 378–387. <https://doi.org/10.1121/1.326851>
- Rojc, M., & Mlakar, I. (2022, May 1). An LSTM-based model for the compression of acoustic inventories for corpus-based text-to-speech synthesis systems. *Computers and Electrical Engineering*, 100, 107942. <https://doi.org/10.1016/j.compeleceng.2022.107942>
- Saleh, H. K., Folkeard, P., Macpherson, E., & Scollie, S. (2020, Jun 8). Adaptation of the connected speech test: Rerecording and passage equivalency. *American Journal of Audiology*, 29(2), 259–264. https://doi.org/10.1044/2019_AJA-19-00052
- Scalart, P., & Filho, J. V. (1996). Speech enhancement based on a priori signal to noise estimation. *ICASSP '96. IEEE International Conference on Acoustics, Speech, and Signal Processing*, 2, 629–632. <https://doi.org/10.1109/ICASSP.1996.541059>
- Schröter, H., Rosenkranz, T., Escalante-B, A. N., & Maier, A. (2022, Aug 15). Low latency speech enhancement for hearing aids using deep filtering. *IEEE/ACM Transactions on Audio, Speech, and Language Processing*, 30, 2716–2728. <https://doi.org/10.1109/TASLP.2022.3198548>
- Scollie, S., Levy, C., Pourmand, N., Abbasalipour, P., Bagatto, M., Richert, F., Moodie, S., Crukley, J., & Parsa, V. (2016, Mar). Fitting noise management signal processing applying the American academy of audiology pediatric amplification guideline: Verification protocols. *Journal of the American Academy of Audiology*, 27(03), 237–251. <https://doi.org/10.3766/jaaa.15060>
- Sharifani, K., & Amini, M. (2023). Machine learning and deep learning: A review of methods and applications. *World Information Technology and Engineering Journal*, 10(07), 3897–3904.
- Taal, C. H., Hendriks, R. C., Heusdens, R., & Jensen, J. (2011, Feb 14). An algorithm for intelligibility prediction of time–frequency weighted noisy speech. *IEEE Transactions on Audio, Speech, and Language Processing*, 19(7), 2125–2136. <https://doi.org/10.1109/TASL.2011.2114881>
- Tay, Y., Dehghani, M., Abnar, S., Chung, H. W., Fedus, W., Rao, J., Narang, S., Tran, V. Q., Yogatama, D., & Metzler, D. (2022, Jul 21). Scaling laws vs model architectures: How does inductive bias influence scaling? arXiv preprint arXiv:2207.10551.
- Thoidis, I., & Goehring, T. (2024). Using deep learning to improve the intelligibility of a target speaker in noisy multi-talker environments for people with normal hearing and hearing loss. *The Journal of the Acoustical Society of America*, 156(1), 706–724. <https://doi.org/10.1121/10.0028007>
- Völker, C., Warzybok, A., & Ernst, S. M. (2015, Dec 16). Comparing bin-aural pre-processing strategies III: Speech intelligibility of normal-hearing and hearing-impaired listeners. *Trends in Hearing*, 19, 2331216515618609. <https://doi.org/10.1177/2331216515618609>
- Weisser, A., & Buchholz, J. M. (2019, Jan 1). Conversational speech levels and signal-to-noise ratios in realistic acoustic conditions. *The Journal of the Acoustical Society of America*, 145(1), 349–360. <https://doi.org/10.1121/1.5087567>
- Wolfe, J., Schafer, E. C., Heldner, B., Mülder, H., Ward, E., & Vincent, B. (2009 Jul). Evaluation of speech recognition in noise with cochlear implants and dynamic FM. *Journal of the American Academy of Audiology*, 20(07), 409–421. <https://doi.org/10.3766/jaaa.20.7.3>
- Wu, Y. H., Stangl, E., Chipara, O., Hasan, S. S., Welhaven, A., & Oleson, J. (2018, Mar 1). Characteristics of real-world signal to noise ratios and speech listening situations of older adults with mild to moderate hearing loss. *Ear and Hearing*, 39(2), 293–304. <https://doi.org/10.1097/AUD.0000000000000486>
- Ye, J., Kobayashi, T., Toyama, N., Tsuda, H., & Murakawa, M. (2018, Aug 13). Acoustic scene classification using efficient summary statistics and multiple spectro-temporal descriptor fusion. *Applied Sciences*, 8(8), 1363. <https://doi.org/10.3390/app8081363>
- Zaar, J., Simonsen, L. B., Behrens, T., Dau, T., & Laugesen, S. (2019). Investigating the relationship between spectro-temporal modulation detection, aided speech perception, and directional noise reduction preference in hearing-impaired listeners. In *Proceedings of the International Symposium on Auditory and Audiological Research*, 7, 181–188.
- Zhao, Y., Wang, D., Johnson, E. M., & Healy, E. W. (2018, Sep 1). A deep learning-based segregation algorithm to increase speech intelligibility for hearing-impaired listeners in reverberant-noisy conditions. *The Journal of the Acoustical Society of America*, 144(3), 1627–1637. <https://doi.org/10.1121/1.5055562>
- Zheng, C., Zhang, H., Liu, W., Luo, X., Li, A., Li, X., & Moore, B. C. (2023). Sixty years of frequency-domain monaural speech enhancement: From traditional to deep learning methods. *Trends in Hearing*, 27, 23312165231209913. <https://doi.org/10.1177/23312165231209913>