

論文 / 著書情報
Article / Book Information

Title	Acoustic feedback cancellation in hearing aids using dual adaptive filtering and gain-controlled probe signal
Author(s)	Muhammad Tahir Akhtar, Felix Albu, Akinori Nishihara
Citation	Biomedical Signal Processing and Control, Vol. 52, pp. 1-13
Pub. date	2019, 7
DOI	http://dx.doi.org/10.1016/j.bspc.2019.03.012
Creative Commons	See next page.
Note	This file is author (final) version.

License



Creative Commons: CC BY-NC-ND

Acoustic Feedback Cancellation in Hearing Aids Using Dual Adaptive Filtering and Gain-Controlled Probe Signal

Muhammad Tahir AKHTAR^{a,*}, Felix ALBU^b, Akinori NISHIHARA^c

^a*Department of Electrical and Computer Engineering, School of Engineering, Nazarbayev University, Astana, Kazakhstan.*

^b*Department of Electronics, Valahia University of Targoviste, 130082 Targoviste, Romania.*

^c*Tokyo Institute of Technology, Ookayama, Meguro-ku, Tokyo 152-8552, Japan.*

Abstract

In this paper, we propose a probe signal-based adaptive filtering method for acoustic feedback cancellation (AFC) in hearing aids. The proposed method consists of two adaptive filters. The first adaptive filter is excited by the receiver (loudspeaker) signal, and uses the microphone signal as its desired response. The first adaptive filter shows a fast convergence speed, however, it may converge to a biased solution at the steady-state because its input and desired response are correlated with each other. The second adaptive filter is excited by an internally generated (uncorrelated) probe signal. The two adaptive filters are adapted using a delay-based normalized least mean square (NLMS) algorithm. A strategy is devised to exchange the coefficients of two adaptive filters such that the both adaptive filters give a good (unbiased) estimate of the acoustic feedback path. Furthermore, we propose to vary the gain of the probe signal, such that a high level probe signal is injected during the transient state, and a low level probe signal is used after the AFC system has converged. The computer simulations demonstrate that the proposed method achieves good modeling accuracy, preserves good speech quality, and maintains high output SNR at the steady-state.

Keywords: Hearing aids, Acoustic feedback, NLMS algorithm, probe signal.

1. Introduction

A typical hearing aid is comprised of an input microphone to pick up the input signal, a processing unit to perform the amplification, noise reduction, and user-dependent frequency-selective processing, and a loudspeaker (receiver). The received signal (from loudspeaker) not only propagates inside the ear, but may also leak-back to the input microphone via the leakage path. This is called ‘acoustic feedback’, and it is a major problem in the hearing aids. The acoustic feedback limits the maximum gain. Furthermore, the presence of feedback results in a closed loop which may initiate oscillations especially when a high gain

*Corresponding Author (Muhammad Tahir AKHTAR)

Email addresses: muhammad.akhtar@nu.edu.kz, akhtar@ieee.org (Muhammad Tahir AKHTAR), felix.albu@valahia.ro (Felix ALBU), aki@cradle.titech.ac.jp (Akinori NISHIHARA)

is applied. This results in the so-called ‘howling’ which is perceived as very annoying whistling and/or screeching sounds. The problem of acoustic feedback in the hearing aids has long been studied [1]–[6], and the most popular approach is to employ the adaptive filtering.

A block diagram of a typical hearing aid, employing the normalized least mean square (NLMS) algorithm-based adaptive filter for acoustic feedback cancellation (AFC), is shown in Fig. 1. Here, $F(z)$, $W(z)$, and $G(z)$ denote the acoustic feedback path, the adaptive filter, and the hearing aid forward path, respectively. Without loss of generality, all transfer functions are considered as finite impulse response (FIR) filters. The signal $x(n)$ picked up by the input microphone comprises two parts: the input signal $s(n)$ which must (ideally) be processed by the hearing aid unit $G(z)$, and the feedback component $y_f(n)$ which is due to a acoustic coupling between the receiver (loudspeaker) and the input microphone. Using $x(n) = s(n) + y_f(n)$ as a desired response and taking the hearing aid signal $y(n)$ as an input (excitation) signal for the AFC filter $W(z)$, the NLMS algorithm [7] to update the coefficients of $W(z)$ is given as

$$\mathbf{w}(n+1) = \mathbf{w}(n) + \frac{\mu e(n) \mathbf{y}(n)}{\|\mathbf{y}(n)\|_2^2 + \epsilon}, \quad (1)$$

where μ is the step-size parameter, $\mathbf{w}(n)$ is the coefficient vector for $W(z)$, $e(n) = x(n) - y_w(n)$ is the error signal, $\mathbf{y}(n)$ is the vector for the output signal $y(n)$, $\|\cdot\|$ denotes the Euclidean norm, and ϵ is a small positive constant to avoid division by zero. Ideally, $W(z)$ attempts to provide a neutralization signal for the acoustic feedback, and hence, the error signal $e(n)$ is used as an input to the hearing aid. The AFC filter $W(z)$ may converge to a biased solution [8] as explained below. In AFC, the forward path results in a closed loop formation of the system which introduces a correlation between the received (loudspeaker) signal $y(n)$ and the desired input signal $x(n)$ causing a biased estimate of the feedback path [9, 10]. This leads to a poor steady-state performance and canceling the desired input signal rather than the feedback signal; hence, introducing the distortion artifacts [11, 12]. Therefore, the scheme of Fig. 1 cannot be used for a continuous AFC in the hearing aids.

A literature review shows that broadly speaking there are two approaches to avoid a biased convergence of the adaptive AFC filter in the hearing aid systems: 1) modify characteristics of the signals present in the hearing aid such that their decorrelation with respect to input signal is improved, and 2) inject internally generated uncorrelated probe signal. In the first class of approaches, the simplest solution is to perform decorrelation by using an appropriate delay either in the cancellation path [1], or in the forward path [13]; however, it degrades the speech quality. Another solution is to filter the error and/or input signal of $W(z)$, through appropriate decorrelation filters, before being used in the update equation of the NLMS algorithm [9], resulting in the so-called Filtered-x adaptive algorithm [14, 15]. The prediction error method (PEM)-based AFC (PEM-AFC) assumes that the input signal can be modeled as an autoregressive process [16, 17, 18]. Once such model is available then the corresponding inverse model can be determined and used to perform the decorrelation. A filter bank-based frequency-domain technique has been investigated

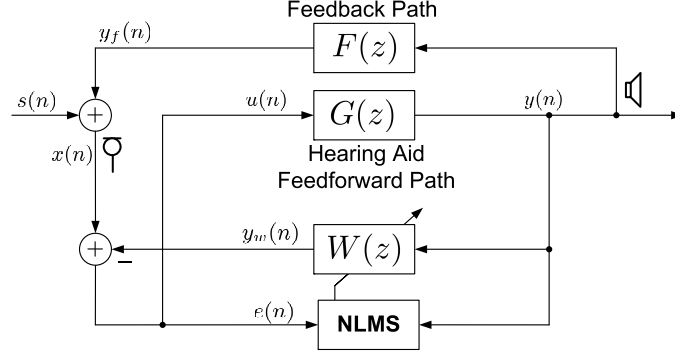


Figure 1: A simplified block diagram of a conventional approach employing the NLMS algorithm for acoustic feedback cancellation in hearing aids.

for AFC [19]. The frequency-domain techniques, however, may require a lot of battery power [2]. For the time-domain continuous AFC, a dual microphone-based solution has been proposed where two microphones are used to pick the input signal and dual adaptive filters are employed to perform AFC [20, 21]. These techniques have obvious physical and computational limitations.

In the probe signal-based approaches, one idea is a noncontinuous adaptation, or an open-loop algorithm in which the hearing aid forward path is broken and a probe signal is injected during the particular intervals, for example, when howling is detected by an appropriate oscillation detector [22]. The ON/OFF switching of the probe signal produces annoying effects to the hearing aid user. A continuous injection of probe signal has been considered, however, either the level of the probe signal must be kept low to have an appreciable signal-to-noise ratio (SNR) [10], or an appropriate probe shaping filter be introduced to perceptually mask the probe signal [23].

In this paper, we consider the probe signal-based adaptation for a continuous AFC in the hearing aids. Essentially, we propose a method comprising two adaptive filters $W_1(z)$ and $W_2(z)$, where $W_1(z)$ is the same as in the conventional approach (Fig. 1), and $W_2(z)$ is excited by an internally generated probe signal. The main features of the proposed method are summarized as follows. 1) A delay is inserted in the path for the probe signal, which allows implementing the delay-based adaptive filtering [24] to track the convergence-status of $W_1(z)$ and $W_2(z)$, 2) a strategy has been developed to transfer the coefficients between the two adaptive filters such that both adaptive filters give a good estimate of the feedback path $F(z)$, 3) the problem of a possible biased convergence is mitigated by freezing the adaptation of $W_1(z)$ once a good solution has been obtained, and finally, 4) a time-varying gain is proposed to control the level of added probe signal: a large value is used at the start-up for a fast convergence, and the gain is reduced to a small value as the system converges thus achieving appreciable SNR at the steady-state. Up to the

best knowledge of authors, the automatic tuning of the probe signal has not been considered in the existing literature on the continuous AFC in the hearing aids.

The rest of the paper is organized as follows. Section 2 reviews the probe signal-based AFC algorithms, and Section 3 briefly describes the PEM-AFC method considered in this paper as a bench mark method for the performance comparison. Section 4 gives details of the proposed method, and Section 5 presents the simulation results. Finally a few conclusions are given in Section 6. A short version of this paper was presented at a conference [25].

Notations: For a sake of convenience of presentation, mixed (time-domain and z -domain) notations have been used as explained here. For an FIR filter having transfer function $H(z)$, its impulse response is denoted as $h(n) = z^{-1}\{H(z)\}$, and the corresponding impulse response coefficient-vector is given as $\mathbf{h}(n) = [h_0(n), h_1(n), \dots, h_{L-1}(n)]^T$ where $h_i(n)$ is the i th sample of $h(n)$. The filtering of an input signal $x(n)$ via $H(z)$ can be represented as $y(n) = h(n) * x(n) = \mathbf{h}^T(n)\mathbf{x}(n)$. Here, z^{-1} and $*$ denote the inverse z -transform and the linear convolution, respectively.

2. Probe Signal-Based AFC

In a classical approach to solve the biased convergence of NLMS algorithm-based AFC, an uncorrelated probe is mixed with the hearing aid signal $y(n)$ [6], [16] (see Fig. 2 (a) without block representing the howling detector). The probe signal $v(n)$ and the hearing aid signal $y(n)$ are together used for continuous adaptation of the AFC filter $W(z)$. Since probe signal would degrade the SNR for the received signal, the level of injected probe signal must be kept low. This, however, results in a slow convergence speed and/or large variance in the estimate [26, 27]. The solution based on howling detector, as shown in Fig. 2(a), would inject the probe noise only when howling is detected and adaptation of the AFC filter is needed. Otherwise, injection of probe noise is stopped. As stated earlier, this appears to be very annoying for the hearing aid users [22].

A somewhat better solution is to use a shaping filter for the probe noise (see Fig. 2(b)). As shown in Fig. 2(b), $v_0(n)$ is an internally generated probe signal which is filtered via the shaping filter $M(z)$ to generate the probe signal $v(n)$ which is mixed with the hearing aid signal $y(n)$ [23]. As shown, this approach employs two copies of the AFC filter $W(z)$ one is adaptive and its coefficients are copied to the otherwise fixed $W(z)$ performing the AFC. It is important to mention that the adaptive AFC filter is excited by only the probe signal $v_0(n)$ in the original formulation. Our experience, however, shows that using both $y(n)$ and $v_0(n)$ improves the convergence as compared with only if $v_0(n)$ is used to excite the AFC filter $W(z)$. Designing an appropriate shaping filter is beyond the scope of the present paper, and the interested reader is referred to [23] and the references there in.

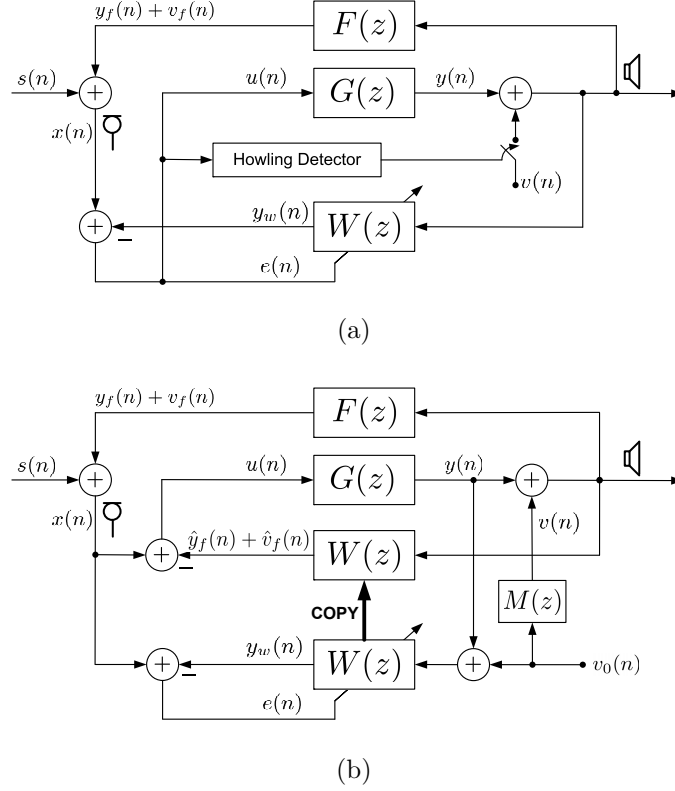


Figure 2: (a) Block diagram for a classical approach for employing a probe noise for an unbiased acoustic feedback cancellation in hearing aids. (b) A modified version of (a) using a shaping filter $M(z)$ for the probe noise.

3. Prediction Error Method (PEM)-Based AFC

The most promising approach to solve the biasing problem in the NLMS-based conventional approach is to employ the PEM-AFC as developed in [16, 17, 18]. Many versions for implementation of PEM-AFC has been considered in the literature, for example, the PEM-based adaptive filtering with row operations (PEM-AFROW) [17, 28], the PEM-based partitioned-block frequency-domain (PEM-PBFD) adaptive filter [29, 30], the PEM-based frequency domain Kalman filter (PEM-FDKF) algorithm [31] among many others. In this paper, we consider a full band time-domain version of PEM-AFC for a sake of a fair comparison with the proposed full band method. Let's consider the block diagram shown in Fig. 3, which shows the PEM-AFC using NLMS algorithm. The basic assumption is that the input signal can be modeled as an AR process ($H(z) = 1/A(z)$) being excited by a white signal. The pre-filtering of the microphone $x(n)$ and the loudspeaker speaker signal $y(n)$ via inverse AR-model (i.e., $A(z)$) would result in whitened signals. This reduces correlation between the input signal $s(n)$ and the rest of signals present in the AFC system.

As shown in Fig. 3, the Levinson-Durbin algorithm [32] is used to perform the AR modeling from the

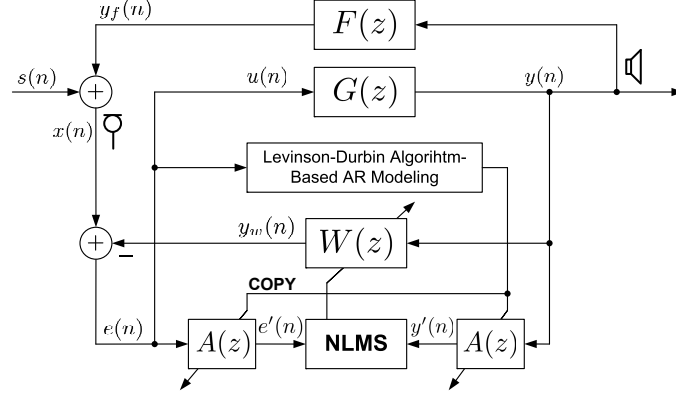


Figure 3: Block diagram of hearing aid employing prediction error method-based acoustic feedback cancellation in hearing aids.

hearing aid input signal $u(n)$ to obtain the coefficients $\mathbf{a}(n) = [1, a_1(n), a_2(n), \dots, a_{L_p}(n)]^T$ of $A(z)$ (here L_p denotes the order of $A(z)$). The coefficients of $A(z)$ are used to obtain the pre-filtered signals $e'(n)$ and $y'(n)$, and finally the coefficients of AFC filter $W(z)$ are updated using the NLMS algorithm as

$$\mathbf{w}(n+1) = \mathbf{w}(n) + \frac{\mu e'(n) \mathbf{y}'(n)}{\|\mathbf{y}'(n)\|_2^2 + \epsilon}. \quad (2)$$

It has been shown [16] that PEM-AFC considerably solves the biased convergence problem of the conventional NLMS-based AFC method. The coefficients of the whitening filter can be obtained offline or updated during the online operation of hearing aid, the latter being considered in this paper. As suggested in [18], for a speech signal at a sampling frequency $F_s = 16$ kHz, the AR model of order $L_p = 10 - 20$ can be updated every 10 – 20 ms.

4. Proposed Method

Figure 4 shows the block diagram of the proposed method for AFC in the hearing aids. As shown, the proposed method consists of two adaptive filters $W_1(z)$ and $W_2(z)$. The first adaptive filter $W_1(z)$ is excited by the receiver (speaker) signal $y(n)$ and is expected to provide a neutralization for the corresponding acoustic feedback component $y_f(n)$. The second adaptive filter $W_2(z)$ is excited by an internally generated uncorrelated probe signal $v(n)$. A time varying gain $\alpha(n)$ (explained later) is used to control the level of $v(n)$ from a white Gaussian noise $v_0(n)$. The second adaptive filter is expected to provide a neutralization signal for the feedback component $v_f(n)$ due to the injected probe signal $v(n)$. Thus, both adaptive filters are adapted to model the characteristics of acoustic feedback path $F(z)$.

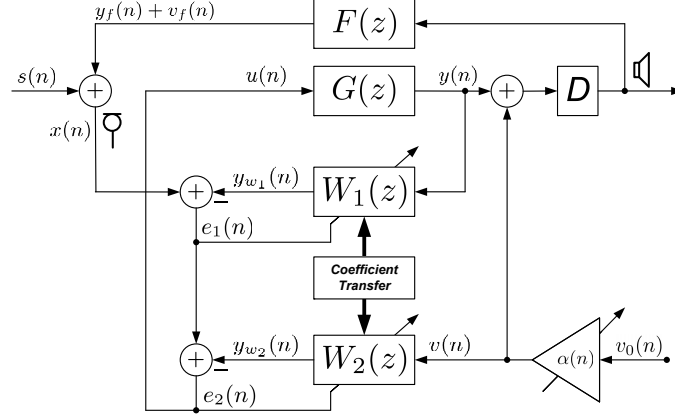


Figure 4: Block diagram of the proposed dual adaptive filtering-based method for continuous acoustic feedback cancellation in hearing aids.

4.1. Adaptation of the AFC Filters

A delay-based technique has largely been applied in the field of acoustic echo cancellation [24]. Here the key idea is to append an ‘appropriate’ delay D in the forward path, and then employ an ‘extended-length’ filter for the system identification. The part of filter employed for modeling the delay is expected to converge to zero. Since the NLMS algorithm spreads the error among the filter coefficients, the norm of extension coefficients can be used as an estimate for the filter mismatch [24]. It is important to note that setting D too low would yield a poor estimator; however, the extension of the adaptive filter implies increased memory and complexity requirements: thus there is a tradeoff situation. Traditionally such type of delay is used to solve the correlation problem in the AFC filter of the hearing aids [13]. In our approach, the objective of the appended delay is two fold: 1) to provide (some) decorrelation, as well as 2) to help designing an efficient strategy for the coefficient transfer between $W_1(z)$ and $W_2(z)$.

Insertion of delay at the output of the hearing aid increases the effective path to be identified by the AFC filters $W_1(z)$ and $W_2(z)$. Thus both $W_1(z)$ and $W_2(z)$ are considered having extended-length coefficient vectors being given as

$$\mathbf{w}_i(n) = \begin{bmatrix} \mathbf{w}_i^D(n) \\ \mathbf{w}_i^F(n) \end{bmatrix}; i = 1, 2, \quad (3)$$

where the first part $\mathbf{w}_i^D(n) = [w_{i,0}(n), w_{i,1}(n), \dots, w_{i,D-1}(n)]^T$ is used to model the appended delay and the second part $\mathbf{w}_i^F(n) = [w_{i,D}(n), w_{i,D+1}(n), \dots, w_{i,D+L-1}(n)]^T$ attempts to model coefficients $\mathbf{f}(n)$ of $F(z)$ which is assumed to be an FIR filter of length L . Both $\mathbf{w}_1^D(n)$ and $\mathbf{w}_2^D(n)$ are initialized with all 1’s and $\mathbf{w}_1^F(n)$ and $\mathbf{w}_2^F(n)$ may be initialized by null vectors. The two adaptive filters are adapted using the

delay-based adaptive algorithms [24] as explained below. In Fig. 4, the input microphone signal $x(n)$ can be expressed as

$$x(n) = s(n) + y_f(n) + v_f(n), \quad (4)$$

where $y_f(n) = f(n) * y(n - D)$ is the acoustic feedback component due to the hearing aid signal $y(n)$ and $v_f(n) = f(n) * v(n - D)$ is the acoustic feedback component due to the probe signal $v(n)$. As stated earlier, the probe signal $v(n)$ is injected for an unbiased convergence of AFC filters. The adaptive filter $W_1(z)$ is excited by $y(n)$ (output of hearing aid forward-path processing unit), and the corresponding error signal for adaptation of $W_1(z)$ is computed as

$$e_1(n) = x(n) - \mathbf{w}_1^T(n) \mathbf{y}(n), \quad (5)$$

where $\mathbf{y}(n) = [\mathbf{y}_D^T(n) \ \mathbf{y}_F^T(n)]^T = [y(n-1), \dots, y(n-D), y(n-D-1), \dots, y(n-D-L)]^T$ is the vector for the hearing aid signal $y(n)$. The coefficients of $W_1(z)$ are updated using the NLMS algorithm as

$$\mathbf{w}_1(n+1) = \mathbf{w}_1(n) + \mu_1(n) \frac{e_1(n) \mathbf{y}(n)}{\|\mathbf{y}(n)\|_2^2 + \epsilon}, \quad (6)$$

where $\mu_1(n)$ is the time-varying step-size being computed as [24]

$$\mu_1(n) = \lfloor \frac{\hat{N}_{D_1}(n)}{P_{e_1}(n) + \epsilon} \rfloor \mu_{\min}, \quad (7)$$

where $\mu_{\min}$ denotes the minimum value of the step-size, $\lfloor \cdot \rfloor$ represents the lower bound operation, and $\hat{N}_{D_1}(n)$ and $P_{e_1}(n)$ are estimated as [24]

$$\hat{N}_{D_1}(n) = \lambda \hat{N}_{D_1}(n-1) + (1-\lambda) \frac{(\|\mathbf{w}_1^D(n)\|_2 \|\mathbf{y}_D(n)\|_2)^2}{D}, \quad (8)$$

and

$$P_{e_1}(n) = \lambda P_{e_1}(n-1) + (1-\lambda) e_1^2(n), \quad (9)$$

respectively, and where λ is the forgetting factor. Using $e_1(n)$ (5) as a desired response and the probe signal $v(n)$ as an input, the corresponding error signal for adaptation of $W_2(z)$ is computed as

$$e_2(n) = e_1(n) - \mathbf{w}_2^T(n) \mathbf{v}(n), \quad (10)$$

where $\mathbf{v}(n) = [\mathbf{v}_D^T(n) \ \mathbf{v}_F^T(n)]^T = [v(n), \dots, v(n-D+1), v(n-D), \dots, v(n-D-L+1)]^T$ is the vector for the probe signal $v(n)$. Similar to (6), the coefficients of $W_2(z)$ are updated as

$$\mathbf{w}_2(n+1) = \mathbf{w}_2(n) + \mu_2(n) \frac{e_2(n) \mathbf{v}(n)}{\|\mathbf{v}(n)\|_2^2 + \epsilon}, \quad (11)$$

where $\mu_2(n)$ is the step-size parameter computed as [24]

$$\mu_2(n) = \lfloor \frac{\hat{N}_{D_2}(n)}{P_{e_2}(n) + \epsilon} \rfloor \mu_{\min}, \quad (12)$$

where $\hat{N}_{D_2}(n)$ and $P_{e_2}(n)$ are estimated as

$$\hat{N}_{D_2}(n) = \lambda \hat{N}_{D_2}(n-1) + (1-\lambda) \frac{(\|\mathbf{w}_2^D(n)\|_2 \|\mathbf{v}_D(n)\|_2)^2}{D}, \quad (13)$$

and

$$P_{e_2}(n) = \lambda P_{e_2}(n-1) + (1-\lambda) e_2^2(n), \quad (14)$$

respectively. As explained below, both adaptive filters $W_1(z)$ and $W_2(z)$ are adapted to give a good estimate of $F(z)$: $y_{w_1}(n) \rightarrow y_f(n)$, $y_{w_2}(n) \rightarrow v_f(n)$, and hence $e_1(n) \rightarrow s(n) + v_f(n) \Rightarrow e_2(n) \rightarrow s(n)$, and hence, $u(n) = e_2(n)$ is used as an input to the hearing aid processing unit $G(z)$.

4.2. Coefficient-Transfer Strategy

In (3), the coefficients corresponding to the appended delay are expected to converge to zero. Therefore, convergence of the two adaptive filters $W_1(z)$ and $W_2(z)$ can be monitored on the basis of the norm of extension coefficients modeling the appended delay as

$$\rho_i(n) = \frac{\|\mathbf{w}_i^D(n)\|_2^2}{D}; i = 1, 2. \quad (15)$$

It is important to note that $W_1(z)$ converges fast, as it is excited by a strong signal (typically $y(n)$ should be an amplified version of the hearing aid input). However, it may eventually converge to a biased solution at the steady-state. To avoid such situation, the following strategy is employed.

1. At the start-up, the convergence of $W_1(z)$ is faster than $W_2(z)$ indicated by $\rho_1(n) < \rho_2(n)$, and hence coefficients of $\mathbf{w}_1^F(n)$ are copied to $\mathbf{w}_2^F(n)$.
2. As $\rho_1(n)$ approaches a certain threshold T_1 , and $W_2(z)$ converges too (indicated by the threshold $\rho_2(n) < T_2$), the adaptation of $W_1(z)$ is frozen, and an estimate of coefficient vector $\mathbf{f}(n)$ (of $F(z)$) is obtained as: $\tilde{\mathbf{f}}(n) = \mathbf{w}_1^F(n)$. Once an estimate of coefficients of $F(z)$ is available, the adaptation of $W_1(z)$ is frozen. This avoid biased convergence of $W_1(z)$ at the steady-state.
3. Now the normalized squared deviation (NSD) for $W_2(z)$ being defined as

$$\text{NSD}_{W_2}(n) = \frac{\|\tilde{\mathbf{f}}(n) - \mathbf{w}_2^F(n)\|^2}{\|\tilde{\mathbf{f}}(n)\|^2}, \quad (16)$$

may be used to monitor any further improvements in the estimation of $F(z)$ by continuous adaptation of $W_2(z)$.

4. The adaptive filter $W_2(z)$ is continuously adapted, and $W_1(z)$ is in fact treated as a piece-wise fixed filter. The parameters $\rho_1(n)$ and $\rho_2(n)$ are computed as in (15). If $\rho_1(n) < \rho_2(n)$ then weights from $\mathbf{w}_1^F(n)$ are copied to $\mathbf{w}_2^F(n)$ and vice versa.

5. If $\text{NSD}_{W_2}(n)$ increases beyond certain predefined threshold T_3 , and the error signal $e_2(n)$ has also diverged (e.g., $P_{e_2}(n) > T_4$), then the acoustic feedback path has changed significantly. The two adaptive filters $W_1(z)$ and $W_2(z)$ are re-initialized to seek a new estimate $\tilde{\mathbf{f}}(n)$ for the changed acoustic feedback path.

4.3. Remarks on Choosing Thresholds

Few remarks on selecting suitable values for the various threshold parameters are given below:

1. As stated earlier, both $\mathbf{w}_1^D(n)$ and $\mathbf{w}_2^D(n)$ can be initialized with all 1's, and are expected to converge to all zeros. Therefore, the parameters $\rho_i(n); i = 1, 2$ in (15) are initialized to 1 at $n = 0$, and $\rho_i(n) \rightarrow 0$ as the AFC system converges. This implies that the threshold T_1 can be selected as a small number close to zero indicating that the overall adaptive AFC filter $W_1(z)$ has converged. Furthermore, the threshold T_2 is used to indicate whether $W_2(z)$ has started converging, and hence can be selected as a number close to but less than unity.
2. It is quite straightforward to understand the selection of the threshold T_3 . The $\text{NSD}_{W_2}(n)$ in (16) computes the misalignment between $\mathbf{w}_2^F(n)$ and $\tilde{\mathbf{f}}(n)$, such that $\text{NSD}_{W_2}(n) = 1$ indicates that $\mathbf{w}_2^F(n) = \mathbf{0}$, $\text{NSD}_{W_2}(n) < 1$ & $\text{NSD}_{W_2}(n) \rightarrow 0$ indicate convergence, and finally $\text{NSD}_{W_2}(n) > 1$ would indicate the divergence. Therefore, $T_3 > 1$ would indicate the divergence of $W_2(z)$ with respect to the estimated feedback path $\tilde{\mathbf{f}}(n)$.
3. T_4 is a threshold for the power of the error signal $e_2(n)$. Since $e_2(n)$ is used as an estimate for the reconstructed speech signal and is input to the hearing aid, the value for the threshold T_4 can be estimated from past values of $P_{e_2}(n)$ computed in (14).

4.4. Time-Varying Gain for Probe Signal

As shown in Fig. 4, the probe signal $v(n)$ is computed from a unit variance white noise $v_0(n)$ as

$$v(n) = \alpha(n)v_0(n), \quad (17)$$

where $\alpha(n)$ is the proposed time-varying gain. Intuitively, we would like to use a high-level probe signal at the start-up (or when there is a change in the acoustic feedback path) so that convergence of $W_2(z)$ is fast. After $W_2(z)$ has converged, the gain $\alpha(n)$ for the probe signal must be reduced to have good output SNR at the steady-state. We have already found a parameter $\rho_2(n)$ which converges to a small value ($\rho_2(n) \rightarrow 0$) as the AFC filter $W_2(z)$ converges. Based on this observation, we propose to compute the time-varying gain $\alpha(n)$ as

$$\alpha(n) = \frac{\rho_2(n)}{\rho_2(n) + C} \quad (18)$$

where C is a positive constant. When $\rho_2(n)$ is large, then $v(n)$ tends to $v_0(n)$. On the other hand, when $\rho_2(n)$ is small, the gain $\alpha(n)$ and hence the probe signal $v(n)$ is small.

Table 1: Summary of the proposed method for continuous AFC in hearing aids.

<i>Parameters:</i> $\epsilon, \lambda, \mu_{\min}, D, L, C, T_1, T_2, T_3, T_4$.
<i>Initialization:</i> $\mathbf{w}_1^D(0) = \mathbf{w}_2^D(0) = \mathbf{1}_{1 \times D}$, $\mathbf{w}_1^F(0) = \mathbf{w}_2^F(0) = \mathbf{0}_{1 \times L}$, FLAG = FALSE , $\alpha(0) = 1$, Generate $v_0(n)$.
while $\{x(n)\}$ available do
$y(n) = g(n) * u(n)$;
Update signal vector $\mathbf{y}(n) = [\mathbf{y}_D^T(n) \mathbf{y}_F^T(n)]^T = [y(n-1), \dots, y(n-D), y(n-D-1), \dots, y(n-D-L)]^T$;
$y_{w1}(n) = \mathbf{w}_1^T(n) \mathbf{y}(n)$;
$e_1(n) = x(n) - y_{w1}(n)$;
$P_{e1}(n) = \lambda P_{e1}(n-1) + (1-\lambda)e_1^2(n)$;
$\hat{N}_{D1}(n) = \lambda \hat{N}_{D1}(n-1) + (1-\lambda) \frac{(\ \mathbf{w}_1^D(n)\ _2 \ \mathbf{y}_D(n)\ _2)^2}{D}$;
$\mu_1(n) = \lfloor \frac{\hat{N}_{D1}(n)}{P_{e1}(n) + \epsilon} \rfloor \mu_{\min}$;
if FLAG == FALSE then update $\mathbf{w}_1(n)$: $\mathbf{w}_1(n+1) = \mathbf{w}_1(n) + \mu_1(n) \frac{e_1(n) \mathbf{y}(n)}{\ \mathbf{y}(n)\ _2^2 + \epsilon}$;
$\rho_1(n) = \frac{\ \mathbf{w}_1^D(n)\ _2^2}{D}$;
$y_{w2}(n) = \mathbf{w}_2^T(n) \mathbf{v}(n)$;
$v(n) = \alpha(n) v_0(n)$;
Update signal vector $\mathbf{v}(n) = [\mathbf{v}_D^T(n) \mathbf{v}_F^T(n)]^T = [v(n), \dots, v(n-D+1), v(n-D), \dots, v(n-D-L+1)]^T$;
$e_2(n) = e_1(n) - y_{w2}(n)$;
$P_{e2}(n) = \lambda P_{e2}(n-1) + (1-\lambda)e_2^2(n)$;
$\hat{N}_{D2}(n) = \lambda \hat{N}_{D2}(n-1) + (1-\lambda) \frac{(\ \mathbf{w}_2^D(n)\ _2 \ \mathbf{v}_D(n)\ _2)^2}{D}$;
$\mu_2(n) = \lfloor \frac{\hat{N}_{D2}(n)}{P_{e2}(n) + \epsilon} \rfloor \mu_{\min}$;
Update $\mathbf{w}_2(n)$: $\mathbf{w}_2(n+1) = \mathbf{w}_2(n) + \mu_2(n) \frac{e_2(n) \mathbf{v}(n)}{\ \mathbf{v}(n)\ _2^2 + \epsilon}$;
$\rho_2(n) = \frac{\ \mathbf{w}_2^D(n)\ _2^2}{D}$;
$\alpha(n) = \frac{\rho_2(n)}{\rho_2(n) + C}$;
$u(n) = e_2(n)$;
if FLAG == FALSE
if $\rho_1(n) < \rho_2(n)$ AND $\rho_2(n) > T_1$ then : $\mathbf{w}_2^F(n) \leftarrow \mathbf{w}_1^F(n)$
elseif $\rho_1(n) < T_1$ AND $\rho_2(n) < T_2$ then : $\tilde{\mathbf{f}}(n) = \mathbf{w}_1^F(n)$, FLAG == TRUE ; endif .
elseif FLAG == TRUE
if $\rho_1(n) < \rho_2(n)$ AND $\rho_2(n) > T_1$ then : $\mathbf{w}_2^F(n) \leftarrow \mathbf{w}_1^F(n)$
elseif $\rho_2(n) < \rho_1(n)$ AND $\rho_2(n) < T_1$ then : $\mathbf{w}_1^F(n) \leftarrow \mathbf{w}_2^F(n)$; endif .
compute $\text{NSD}_{W_2}(n) = \frac{\ \tilde{\mathbf{f}}(n) - \mathbf{w}_2^F(n)\ ^2}{\ \tilde{\mathbf{f}}(n)\ ^2}$;
if $\text{NSD}_{W_2}(n) > T_3$ AND $P_{e2}(n) > T_4$ AND $\rho_2(n) > 1$ then :
Re-initialize AFC filters $\mathbf{w}_1(n)$ and $\mathbf{w}_2(n)$, and set FLAG = FALSE ; endif .
endif
end while

4.5. Computational Complexity

In this section, details of the computational complexity are given in terms of the computations required per iteration of the adaptive filtering to perform AFC in the proposed method in comparison with the PEM-AFC method. It is straightforward to carry out such analysis for the methods considered in this paper, hence, details are omitted and only final expressions are given. The computational summary of the proposed method is given in Table 1. The proposed method requires $12L + 16$ multiplications, $12L + 1$ additions, and 7 divisions per iteration of execution. On the other hand, the NLMS algorithm based PEM-AFC method (Fig. 2) requires $3L + 2L_p + 1$ multiplications, $3L + 2L_p - 2$ additions, and one division per iteration. In addition, PEM-AFC method would require executing Levinson-Durbin algorithm to update coefficients of whitening filter $A(z)$ at regular intervals. We may, therefore, conclude that the computational complexity of the proposed method is comparable to that of PEM-AFC approach.

5. Simulation Results

In this section, we present results of the computer simulations, where we consider the following methods for the performance comparison:

1. The NLMS algorithm-based conventional method.
2. The NLMS algorithm with probe shaping as described in Section 2. The filter to shape the probe noise is adopted from [23], and is given as $M(z) = 1/[1 - 2 \times 0.92 \cos(\frac{200 \times 2\pi}{15750})z^{-1} + 0.92^2 z^{-2}]$.
3. The PEM-AFC method using the NLMS algorithm where PEM-based filtering is used to modify the reference and error signal used in the adaptation of AFC filter. The MATLAB function `aryule`, which implements Levinson-Durbin, is used to update coefficients of (AR modeling-based) $L_p = 16$ order pre-filter $A(z)$ every 10 ms.
4. The proposed method as described in Section III.

The acoustic feedback path $F(z)$ is modeled as an FIR filter of length $L = 64$. The corresponding data is taken from [33] where the feedback path measurements have been performed over human subjects, and under various conditions, viz., closed mouth, open mouth, wide open mouth, and telephone close to the ear [33]. For further details on the experiments and measurements, the interested reader is referred to [33]. Fig. 5 shows the impulse and magnitude response characteristics of the acoustic feedback path used for simulations carried out in this paper. The sampling frequency is $F_s = 16$ kHz, and all adaptive filters are assumed to be FIR filters of length 64. For the performance comparison, the following performance measures have been employed.

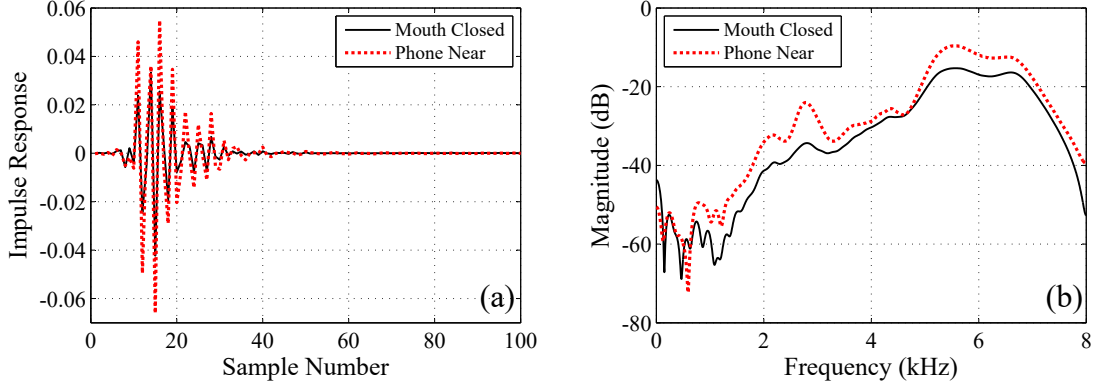


Figure 5: The impulse response (a) and magnitude response (b) characteristics of electro-acoustic feedback path $F(z)$ used in the computer simulations.

- *Normalized Squared Deviation (NSD)*: The misalignment of AFC method with respect to the true feedback path $F(z)$ is computed as NSD as follows

$$\text{NSD}_{W_1}(n)(dB) = 10 \log \left\{ \frac{\|\mathbf{f}(n) - \mathbf{w}_1^F(n)\|^2}{\|\mathbf{f}(n)\|^2} \right\}, \quad (19)$$

- *Maximum Stable Gain (MSG)*: The MSG is computed as [21]

$$\text{MSG}(dB) = 10 \log \left\{ \max_{\omega} \|F(\omega) - W_1^F(\omega)\|^2 \right\}, \quad (20)$$

where $F(\omega)$ and $W_1^F(\omega)$ denotes Fourier transform of feedback path coefficients $\mathbf{f}(n)$ and the AFC filter (corresponding) coefficients $\mathbf{w}_1^F(n)$, respectively. As noted above, the MSG is determined by the frequency where the mismatch between the actual and the estimated path is greatest. However, the system will only be unstable when the phase at that frequency equals a multiple of 2π [21].

5.1. Case 1: Speech Signals

In this case study, the experiments have been performed for the set of speech signals shown in Fig. 6. The forward path representing the hearing-aid processing unit, is assumed to be given as $G(z) = Kz^{-\Delta}$, where K and Δ , respectively, represent the gain and delay of the system. The experiments have been performed for $K = 10$ (a low gain scenario), and $K = 30$ (a high gain scenario). The delay is selected as $\Delta = 80$ which corresponds to 5 ms. The simulation parameters have been determined experimentally (by trial-and-error) for fast and stable convergence, and have been adjusted as follows:

- NLMS algorithm-based conventional and PEM-AFC methods: $\mu = 1 \times 10^{-3}$, $\epsilon = 1 \times 10^{-4}$.
- NLMS algorithm with probe shaping: $\mu = 1 \times 10^{-3}$, $\epsilon = 1 \times 10^{-4}$, $\text{SNR}_{\text{probe}} = \sigma_v^2 / \sigma_s^2 = -20$ dB.

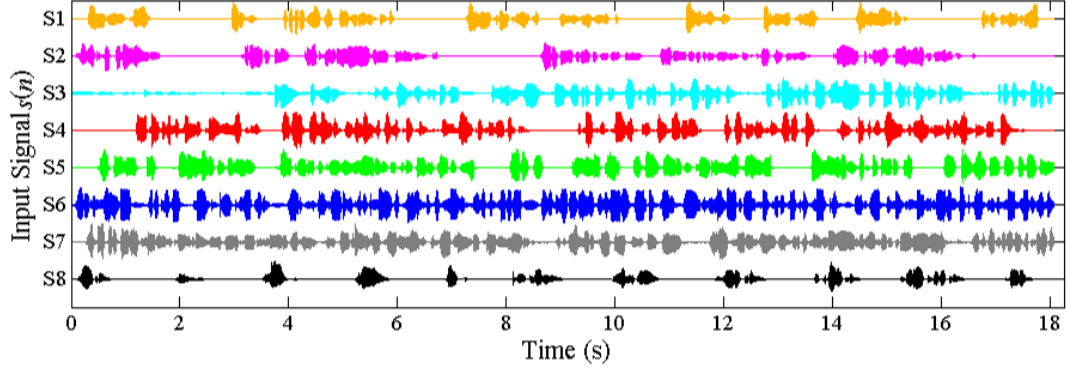


Figure 6: Plots for various speech signals used in the computer simulations in Case 1.

- Proposed method: $\lambda = 0.97$, $D = 40$, $\mu_{\min} = 1 \times 10^{-6}$, $T_1 = 1 \times 10^{-3}$, $T_2 = T_3 = 1$, $T_4 = 10$, $\text{SNR}_{\text{probe}} = \sigma_v^2 / \sigma_s^2 = 0$ dB, $C = 1.5$. It is worth to mention that the delay $D = 40$ corresponds to 2.5 ms for a sampling frequency of $F_s = 16$ kHz, and hence, the total processing delay is 7.5 ms in the proposed method. This corresponds to an acceptable processing delay in the hearing aids [34].

The simulation results for the hearing aid model with $K = 10$ for signals S1 to S8 are presented in Figs. 7-10. Figure 7 shows the error in the reconstruction of the desired signal being computed as $\Delta S(n) = |s(n) - u(n)|$ at the input of the hearing aid for the speech signal S6. It is obvious that for a perfect reconstruction, we must have $\Delta S(n) \rightarrow 0$. From Fig. 7, we see that the proposed method gives a fast convergence speed and small steady-state error in reproducing the desired signal at the input of hearing aid processing unit $G(z)$.

Fig. 8 shows the impulse response and magnitude response of AFC filters $W(z)$ ($W_1(z)$ in the proposed method) at the steady-state in comparison with the true feedback path $F(z)$. The results are presented for the speech signal S6 as shown in Fig. 6. We observe that the PEM-AFC algorithm and the proposed method algorithm clearly outperform the rest of the algorithms. Furthermore, the proposed method gives the best performance among the algorithms studied in this paper, whether observed in the time or the frequency domain (see Figs. 8 (g) and (f)).

Fig. 9(a) shows the performance comparison for NSD (19), where each curve is obtained by averaging for all speech signals S1 to S8. The proposed method do not update $W_1(z)$ once a good solution is obtained. This avoids any fluctuations due to the non-stationarity of speech signal. We observe that the proposed method gives a very fast convergence speed and a good steady-state performance as compared with the other methods. In order to highlight the performance comparison for the rest of algorithms, two small panels are shown in the same Fig. 9(a). The convergence speed of the probe shaping NLMS algorithm is indeed faster

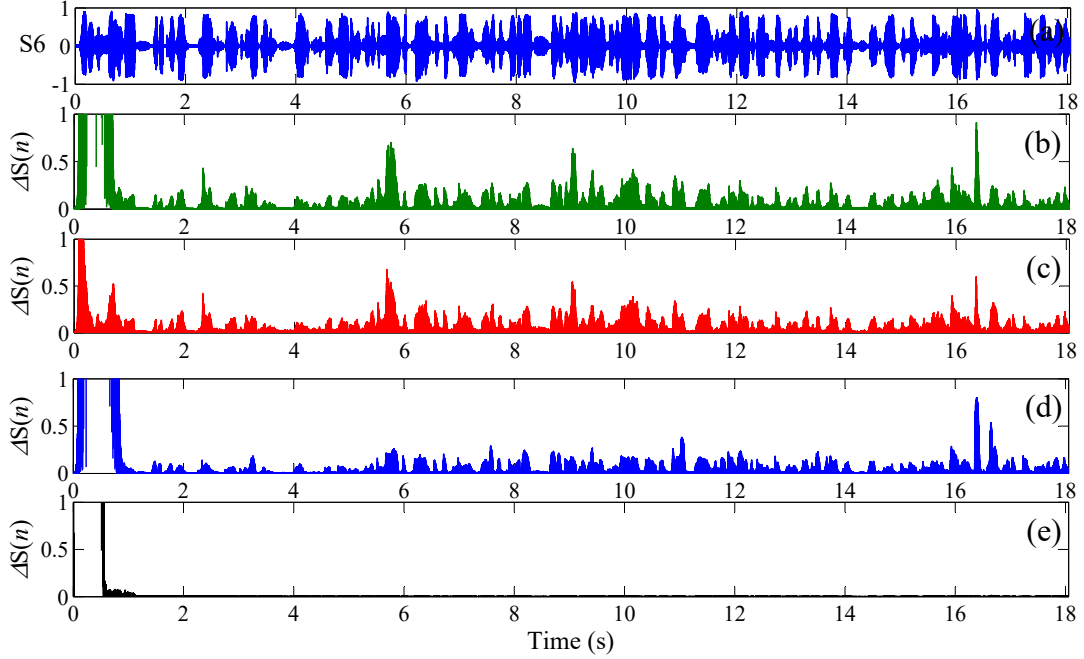


Figure 7: Simulation results in Case 1. (a) The input signal $s(n)$ (S6 in Fig. 6), and (b)-(e) the corresponding error in the reconstruction of the input signal: $\Delta S(n) = |s(n) - u(n)|$, respectively, for the NLMS algorithm, probe shaping NLMS, PEM-AFC, and the proposed method.

than the rest of algorithms. On the other hand, the PEM-AFC method gives a better NSD performance at the steady-state as compared with the other algorithms. The corresponding curves for MSG (20) are shown in Fig. 9(b), where we observe that the proposed method gives largest MSG as compared with the other methods considered in this paper.

The curves for variation of parameters $\rho_1(n)$ and $\rho_2(n)$ (as defined in (15)) are shown for the proposed method in Fig. 10(a). As expected, $W_1(z)$ converges very fast indicated by the fast convergence of $\rho_1(n)$. As $\rho_1(n)$ drops below (the pre-decided threshold) T_1 and $W_2(z)$ converges too (indicated by T_2 on $\rho_2(n)$), the adaptation of $W_1(z)$ is stopped to avoid its biased convergence. In the proposed method, $W_1(z)$ is the main AFC filter and $W_2(z)$ acts mainly as a supporting filter. The convergence speed of $W_1(z)$ is very fast, therefore, $\rho_1(n)$ drops to very low level. As soon as $W_2(z)$ converges, its input $v(n)$ is decreased to a very low level. In the presence of $s(n)$ (acting as a strong disturbance in the desired response of $W_2(z)$ as compared with the low level input $v(n)$), $\rho_2(n)$ settles to a level which is not as low as achieved by $\rho_1(n)$. Now the role of $W_2(z)$ is mainly to monitor the status of AFC system, and re-initialize the adaptation if required for example for any change in the acoustic environment. In the absence of input speech signal (for example during the long silence periods when nobody is talking to the hearing aid user; however, the hearing aid is

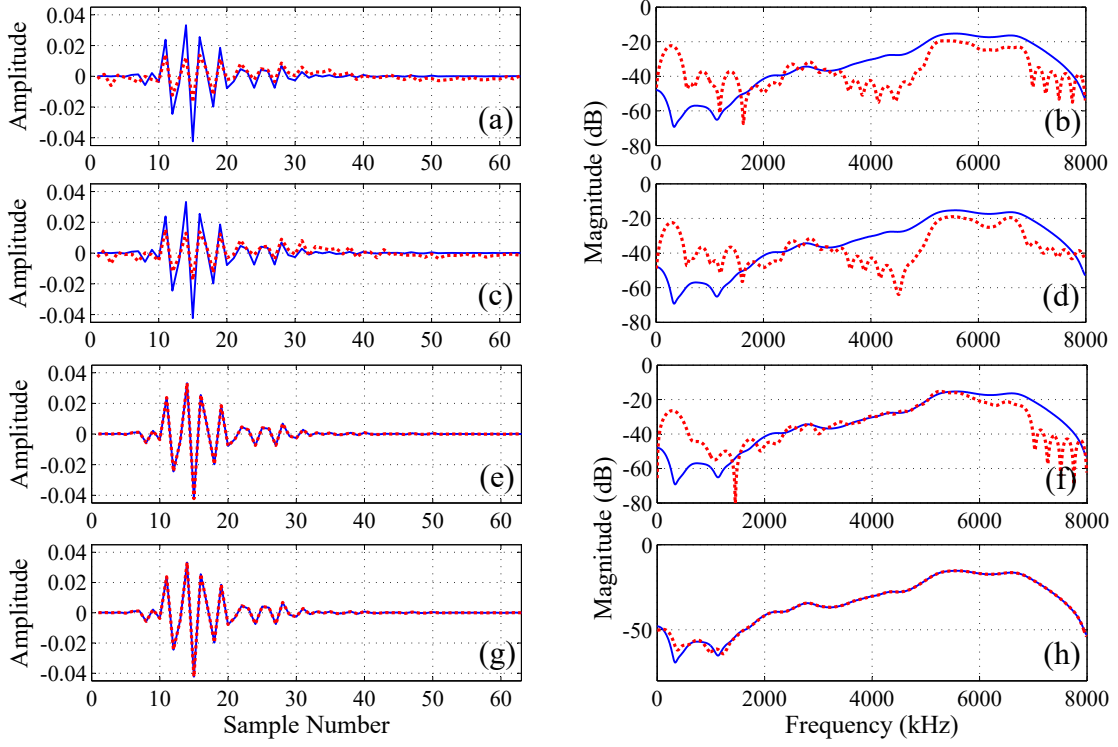


Figure 8: Simulation results in Case 1. The steady-state impulse response (left panel) and magnitude response (right panel) characteristics of the acoustic feedback path model (dotted line) by various algorithms in comparison with the original feedback path (solid line). The presented results have been obtained for the input signal S6 by employing the NLMS algorithm ((a) and (b)), probe shaping NLMS ((c) and (d)), PEM-AFC ((e) and (f)), and the proposed method ((g) and (h)).

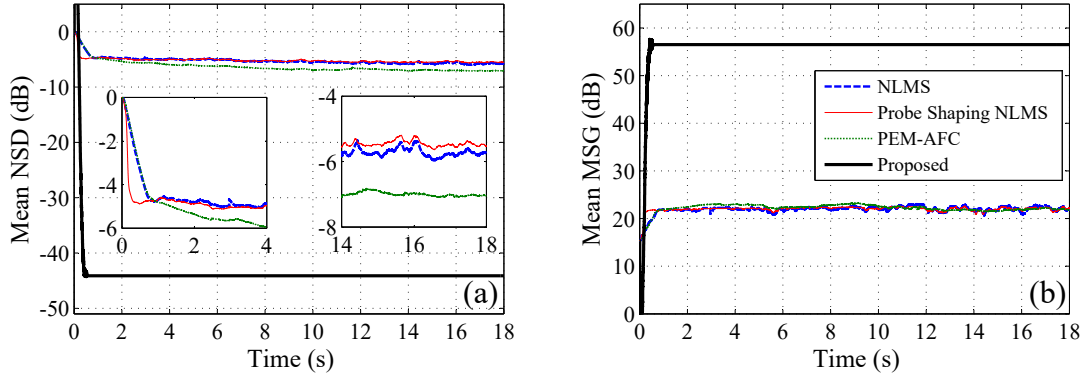


Figure 9: Simulation results in Case 1. (a) Plot of normalized squared deviation (NSD) and (b) maximum stable gain (MSG) averaged over the speech signals S1 to S8 shown in Fig. 6.

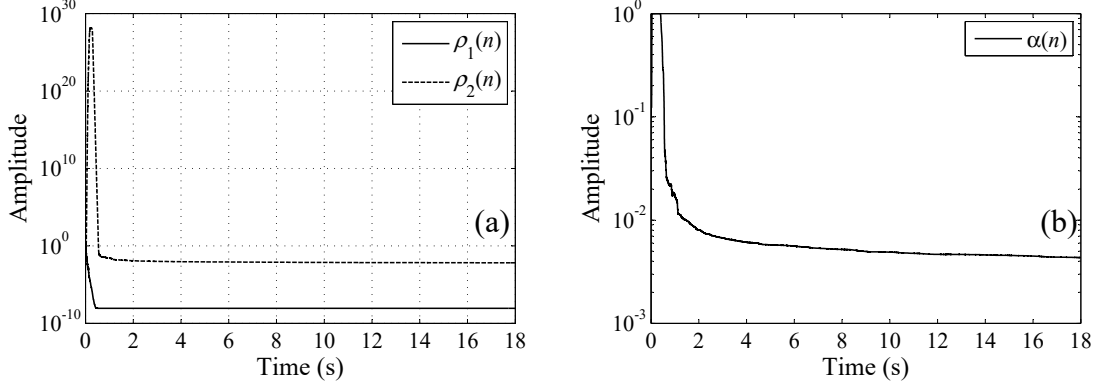


Figure 10: Simulation results in Case 1. (a) Variation of parameters $\rho_1(n)$ and $\rho_2(n)$ (as defined in (15)), and (b) time varying gain $\alpha(n)$ (as defined in (18)) in the proposed method (averaged over the speech signals S1 to S8).

still ON), $W_2(z)$ may further (slowly) adapt to give an improved (fine-tuned) estimate of the feedback path. Such situation, however, is not considered in the results presented in this paper. Figure 10(b) shows the variation of time varying gain $\alpha(n)$ (as defined in (18)) for the probe signal $v(n)$ in the proposed method. We observe that at the start-up a large gain is used for the probe signal resulting in a fast convergence of the AFC filter. At the steady-state, the gain (for probe signal) reduces to a very small value which in turn improves the SNR at the output of hearing aid (as explained later).

The above experiment is repeated for a hearing aid model with $K = 30$ (a high amplification scenario). The results are not shown for the sake of duplication. The quantitative analysis for various methods is carried out for the following performance measures:

- *Hearing aid speech quality index (HASQI)*: It is an index especially developed for the speech signals processed in hearing aids. HASQI is developed to predict the speech quality variations for noise, nonlinear distortions, the degradations caused by frequency compression, noise suppression, speech vocoding, acoustic feedback and feedback cancellation, and speech combined with modulated noise [35, 36]. The maximum score of 1.0 is for a clean signal with no degradation.
- *Hearing aid audio quality index (HAAQI)*: It is an index developed to predict the audio quality (especially for music signals) in the hearing aids, and employs the HASQI auditory model with different parameters fitted to the data of an extensive music-quality rating experiment [37]. The maximum score of 1.0 is for a clean signal with no degradation.
- *Perceptual Evaluation of Speech Quality (PESQ)*: It is an ITU-T standard to evaluate the quality of the speech signals [38]. The maximum score of 4.5 is for a clean signal with no degradation.

- *Signal to Distortion Ratio (SDR)*: It is based on the Hilbert transform and measure levels of the nonlinear distortion in the processed signal in comparison with the original signal [39, 40, 41]. In the simulation model (see Figs. 1-4), the input signal $s(n)$ and the reconstructed signal $u(n)$ are taken as reference and test signals, respectively.
- *Mutual Information (MI)*: It is a non-parametric measure of relevance between two random variables z_1 and z_2 , and can be interpreted using the Kullback-Leibler divergence as [42]

$$\text{MI} = \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} f(z_1, z_2) \log \left(\frac{f(z_1, z_2)}{f(z_1)f(z_2)} \right) dz_1 dz_2, \quad (21)$$

where $f(z_1, z_2)$ denotes the joint probability distribution function of z_1 and z_2 , and $f(z_1)$ and $f(z_2)$ denote their respective marginal probability distribution functions (PDFs). MI is always non-negative and zero if and only if the two random variables are strictly independent. We compute MI between input $s(n)$ and the reconstructed signal $u(n)$, with the understanding that $(\text{MI} > 0)$ indicates close relevance between these signals, and larger the MI better $u(n)$ resembles the input $s(n)$.

- *Normalized Mean Squared Error (NMSE)*: It is computed (in dB) as

$$\text{NMSE} = 10 \log \left\{ \frac{\frac{1}{J} \sum_j [s(j) - u(j)]^2}{\frac{1}{J} \sum_j [s(j)]^2} \right\}, \quad (22)$$

where J denotes the total number of samples used in the computation. Being a ratio of two similar quantities, NMSE is a unit-less quantity, and $\text{NMSE} \rightarrow -\infty$ shows that the corresponding signal is reconstructed with the minimum error.

- *Averaged MSG*: This is obtained by averaging the steady-state value of MSG across all speech signals considered in the experiment.
- *Output SNR*: For the NLMS algorithm with probe shaping and the proposed methods, the output SNR may be computed as $\text{SNR}_{\text{out}} = 10 \log \{\sigma_y^2 / \sigma_v^2\}$, where σ_y^2 and σ_v^2 denote variances of output signal $y(n)$ and probe signal $v(n)$, respectively.

It is important to note in the above-mentioned objective measures, the most are not designed to evaluate the sound quality against feedback and AFC artifacts. Furthermore, some of them are well suited to access the perceptual aspects of speech but not music signals [43]. Nevertheless, we employed the above-mentioned performance measures to have quantitative assessment of various method in the same settings, and hence, the performance comparison is still valid. Table 2 summarizes the quantitative results for the above-mentioned performance measure averaged over all speech signals S1-S8 (from mid sample to the last value). It can be observed that the proposed method gives best performance among the methods considered in this paper. The speech quality is severely degraded for the NLMS, probe shaping NLMS, and PEM-AFC algorithms,

Table 2: Quantitative assessment of various methods for speech signals S1 to S8 in Case 1.

		$K = 10$				$K = 30$			
		NLMS	NLMS-Probe	PEM-AFC	Proposed	NLMS	NLMS-Probe	PEM-AFC	Proposed
HASQI	Mean	0.8310	0.8345	0.8157	0.9963	0.8268	0.8032	0.7053	0.9914
	SD	0.0618	0.0681	0.2209	0.0015	0.0577	0.0666	0.3275	0.0021
	Median	0.8574	0.8396	0.9013	0.9962	0.8525	0.8166	0.8644	0.9906
HAAQI	Mean	0.7860	0.7202	0.8603	0.9970	0.7644	0.6932	0.8093	0.9935
	SD	0.0601	0.0863	0.0953	0.0005	0.0557	0.0643	0.1456	0.0006
	Median	0.7697	0.7282	0.8882	0.9970	0.7530	0.6815	0.8716	0.9937
PESQ	Mean	3.6141	3.4652	3.8348	4.4282	3.6030	3.3309	3.6030	4.3582
	SD	0.4002	0.4422	0.9700	0.0825	0.3949	0.4516	1.0184	0.1929
	Median	3.7205	3.6306	4.1413	4.4832	3.7198	3.3423	4.0135	4.4481
SDR	Mean	9.1338	9.6438	15.9767	46.8282	9.0748	7.8464	15.4676	45.8336
	SD	2.4931	2.5259	3.7300	9.8639	2.5881	2.1360	4.1433	9.4086
	Median	9.0777	9.9817	15.7640	51.0084	9.0378	7.5642	15.6999	49.6905
MI	Mean	0.9753	0.9981	1.1454	2.4927	0.9638	0.8887	0.9890	2.2596
	SD	0.1813	0.1832	0.2793	0.4903	0.1763	0.1455	0.4816	0.4363
	Median	0.9658	0.9777	1.1646	2.5741	0.9562	0.9202	1.1325	2.3301
NMSE	Mean	-7.1150	-8.3918	-8.4073	-34.6052	-7.0086	-7.0152	-1.6565	-27.7552
	SD	2.1527	2.0058	7.0388	4.0285	2.2653	1.8719	16.7417	2.2177
	Median	-7.2069	-8.9354	-10.9836	-34.5468	-7.3262	-7.0277	-9.9002	-27.4118
MSG	Mean	22.1291	22.0364	22.0751	56.5088	33.6283	31.3236	33.9974	63.1463
	SD	0.8253	1.0665	1.2147	6.4870	1.1400	0.6648	1.8542	3.1898
	Median	22.0885	21.8852	21.6992	57.0987	34.0613	31.3936	34.8070	62.9620
SNR _{out}	Mean	–	15.3477	–	66.9031	–	24.2555	–	76.3205
	SD	–	1.0420	–	1.2806	–	1.0980	–	1.5075
	Median	–	15.4046	–	66.5453	–	24.2043	–	76.1260

as depicted by their low SDR values. The proposed method gives the best performance as compared with the rest of methods considered in this paper. Especially, the values for HASQI, HAAQI, PESQ, MSG, and output SNR have been substantially improved as compared with the other methods. As described earlier, the proposed method employs a time varying gain $\alpha(n)$ for the probe signal $v(n)$, such that the probe signal automatically reduces to a very low level at the steady-state as the AFC system converges. This substantially improves the output SNR, and in fact the probe signal is not audible at the steady-state. **It is worth to mention that steady-state output SNR for the probe shaping NLMS algorithm is very low. This is due to the reason that (shaped) probe noise appears as a background noise present all the time.**

5.2. Case 2: Signals with Tonal Characteristics

In this case study, we consider signals with strong tonal characteristics. See Table 3 for description of signals and the corresponding plots are given in Fig. 11. These signals are hereby used to evaluate the

Table 3: Description of input signals considered in Case 2. (The corresponding wave plots are shown in Fig. 11.)

S9	Ambulance siren
S10	Strong tone with some speech component
S11	Impulsive signal with periodic tones
S12	Police siren with ambulance siren in the background
S13	Impulsive signal with quasi-periodic tones
S14	Typewriter sound
S15	Barking puppy (very shrilling sound)

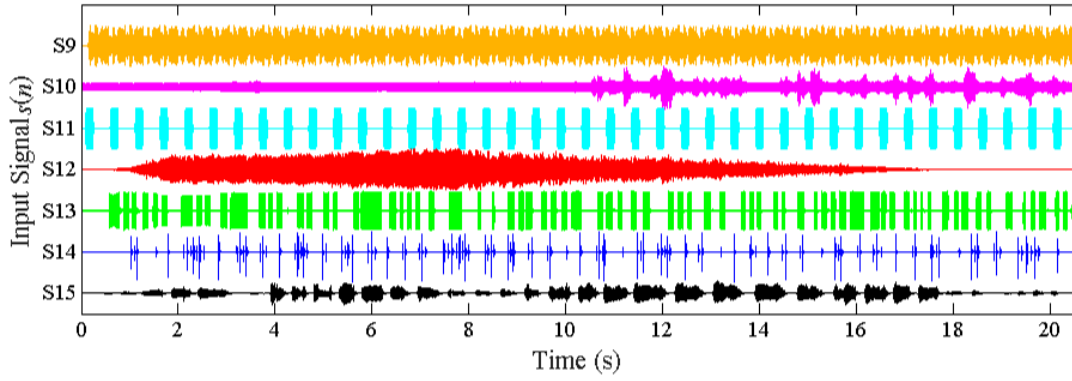


Figure 11: Plots for various speech signals (with strong audio content) used in the computer simulations in Case 2.

performance of hearing aid for entrainment. Entrainment is typically described as feedback after cessation of the sound, additional tones, warbling, or echoes [44]. It is a common artifact in hearing aids which occurs when the feedback cancellation algorithm erroneously attempts to cancel a tonal input to the hearing aid [45]. The simulation parameters are adjusted to the same values as in Case 1, and the corresponding simulation results for hearing aid model with $K = 10$ are presented in Figs. 12-15.

Fig. 12 shows error in the reconstruction of the desired signal at the input of the hearing aid for the input signal S10. As mentioned in Table 3, S10 comprises a strong tone present all the time with some speech content appearing at about 10 s. This mimics the situation when a buzzer alarm is constantly ringing along with some announcement, for example, to evacuate a building in an emergency. The performance of NLMS algorithm is very poor. In fact, we hear a lot of musical noise and tones even not present in the input signal. The probe shaping can improve the convergence speed, and gives somewhat better performance in replicating the (input) signal for exciting the hearing aid. As noted in [18], the basic assumption that the input signal can be modeled as an AR process being excited by a white signal is not valid for such strongly correlated signals. Therefore, PEM-AFC method also suffers from some musical noise, though its performance is better

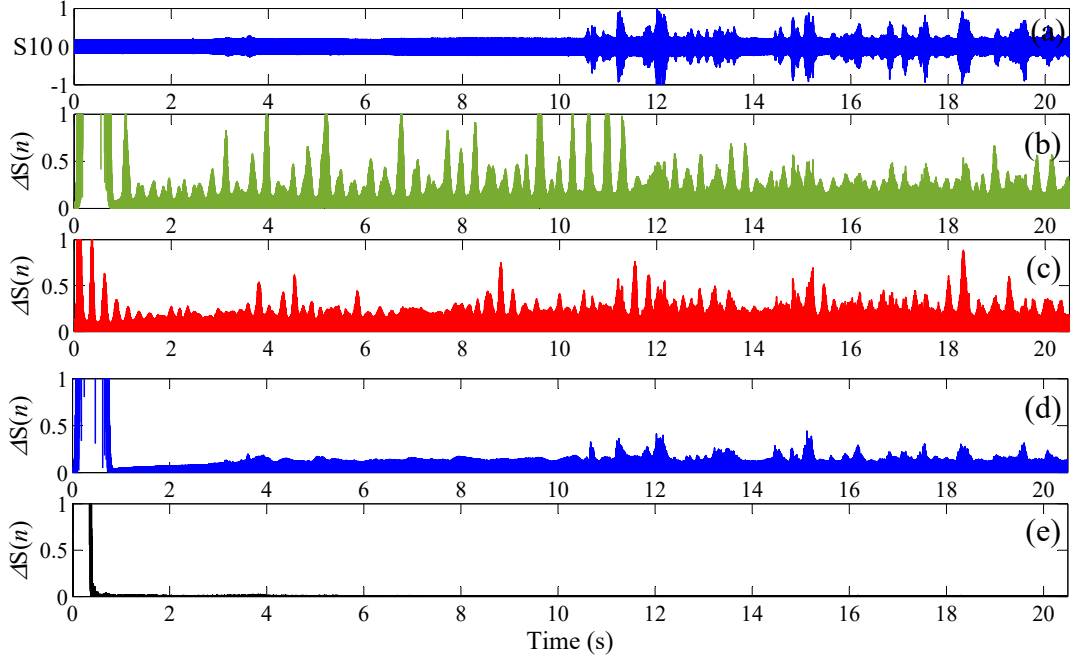


Figure 12: Simulation results in Case 2. (a) The input signal $s(n)$ (S10 in Fig. 11), and (b)-(e) the corresponding error in the reconstruction of the input signal: $\Delta S(n) = |s(n) - u(n)|$, respectively, for the NLMS algorithm, probe shaping NLMS, PEM-AFC, and the proposed method.

than the NLMS and probe shaping NLMS algorithms. On the other hand, the proposed method gives very fast convergence speed and smallest error in generating the hearing aid input signal $u(n)$.

The corresponding spectrograms of the input signal $s(n)$ and hearing aid input signal $u(n)$ for various methods, are shown in Fig. 13 for the signal S10. As stated earlier, this signal comprises a strong tonal component present all the time and includes a speech portion as well. It can be noticed from Fig. 13(a) that the conventional NLMS-based approach suffers from the entrainment artifacts, and many frequencies not present in the original signal can be seen in the spectrogram shown in Fig. 13(a). A lot of musical noise can be heard in the case of conventional method. The probe shaping NLMS algorithm also suffers from the entrainment artifacts, though not as severe as in the case of conventional method using NLMS algorithm. The PEM-AFC method gives better performance than the NLMS and probe shaping NLMS algorithms. Furthermore, the proposed method produces no such musical noise, gives best performance in terms of retaining tonal as well as speech contents of the input signal.

The mean NSD and MSG curves, averaged over signals S9 to S15, are shown in Fig. 14, and the corresponding curves for convergence of the proposed method studied in terms of parameters $\rho_1(n)$, $\rho_2(n)$, and $\alpha(n)$ are shown in Fig. 15. It can be noticed that the proposed method clearly outperforms the rest of the

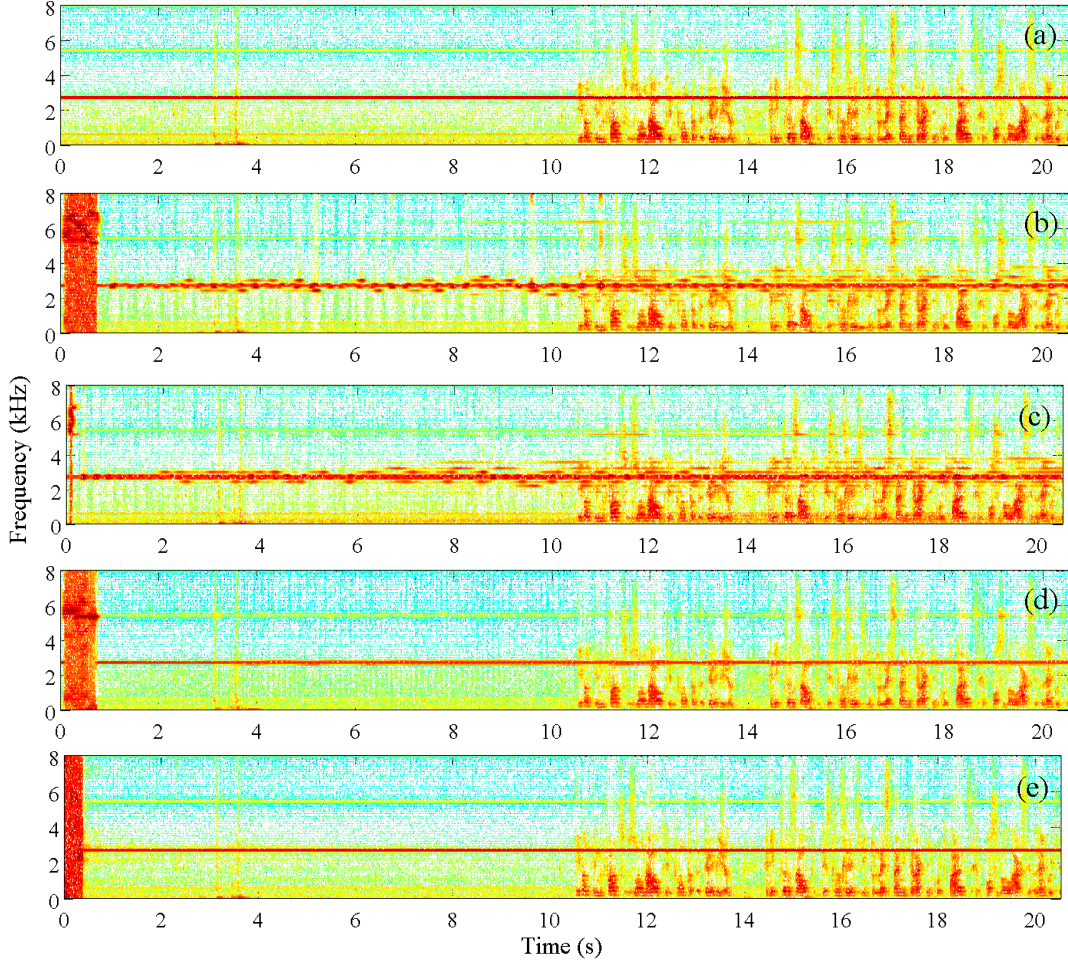


Figure 13: Simulation results in Case 2. (a) Spectrogram of input signal $s(n)$ (S10). (b)-(e) Spectrograms for hearing aid (reconstructed) input signal $u(n)$ in the NLMS algorithm, probe shaping NLMS, PEM-AFC method, and proposed method, respectively.

methods considered in this paper. The experiments for $K = 30$ show similar performance behavior (results not shown). Table 4 summarizes the quantitative results for various performance measures. It is evident that the proposed method outperforms the other methods with a margin in all performance measures. As stated earlier, the input signals considered in this case have strong correlation characteristics and appear as a very tough challenge for the existing algorithms considered in this paper.

5.3. Case 3: Sudden Change in Acoustic Feedback Path

In this case study, a sudden change in the acoustic feedback path $F(z)$ is considered. This situation may arise in practice when the hearing aid user brings, for example, mobile phone near to his/her ear. It results

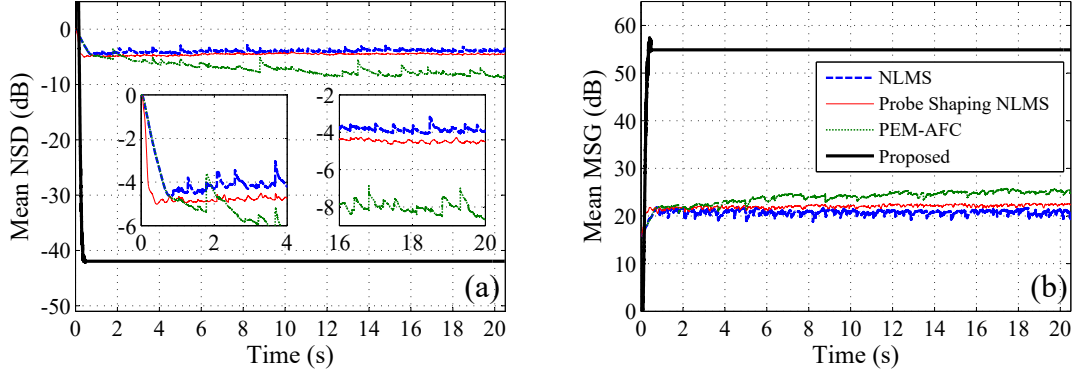


Figure 14: Simulation results in Case 2. (a) Plot of normalized squared deviation (NSD) and (b) maximum stable gain (MSG) averaged over the speech signals S9 to S15 shown in Fig. 11.

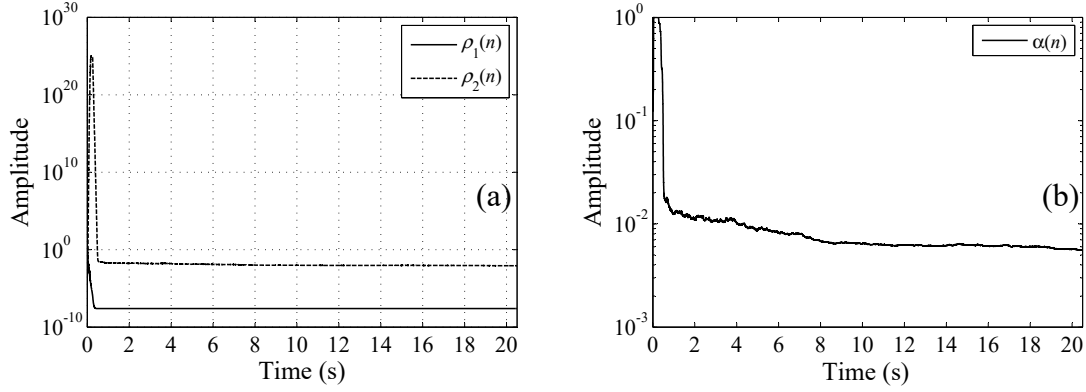


Figure 15: Simulation results in Case 2. (a) Variation of parameters $\rho_1(n)$ and $\rho_2(n)$ (as defined in (15)), and (b) time varying gain $\alpha(n)$ (as defined in (18)) in the proposed method (averaged over the speech signals S9-S15).

in the increase of amplitude of the impulse response of the feedback path [33]. At the startup, the acoustic path is the same as considered in the previous case. At the middle of the simulation, the acoustic feedback is suddenly changed to a new one [which is in fact measured during such an above-mentioned situation](#). The characteristics of the changed acoustic feedback path are shown by dotted curves in Fig. 5. The simulation parameters are adjusted to the same values as found in Case 1, and the hearing aid gain is adjusted to $K = 10$. The corresponding simulation results are shown in Fig. 16. We observe that the proposed method keeps good performance before as well as after the sudden change in the acoustic feedback path. The gain for probe signal rises to a large level when acoustic path changes, and automatically reduces to a low level as the AFC system converges.

Table 4: Quantitative assessment of various methods for signals S9 to S15 in Case 2.

		$K = 10$				$K = 30$			
		NLMS	NLMS-Probe	PEM-AFC	Proposed	NLMS	NLMS-Probe	PEM-AFC	Proposed
HASQI	Mean	0.6112	0.6191	0.7127	0.9933	0.5906	0.6219	0.6808	0.9869
	SD	0.3160	0.3062	0.3692	0.0071	0.3093	0.3067	0.3569	0.0116
	Median	0.8574	0.7049	0.9152	0.9970	0.5051	0.5616	0.8055	0.9899
HAAQI	Mean	0.5117	0.4995	0.6579	0.9852	0.4994	0.5371	0.6196	0.9662
	SD	0.2406	0.2170	0.3361	0.0175	0.2397	0.2516	0.3173	0.0394
	Median	0.5297	0.6413	0.7787	0.9946	0.7530	0.5537	0.6370	0.9821
PESQ	Mean	3.0080	3.0690	3.5406	4.4209	2.9467	3.1110	3.4912	4.3868
	SD	0.9651	0.8438	1.2420	0.0710	0.9252	0.9640	1.1650	0.1438
	Median	2.7356	3.3827	4.1115	4.4128	2.7319	2.7847	3.7478	4.4490
SDR	Mean	11.3318	19.1240	17.2929	49.3712	11.4506	15.7399	15.9143	47.7864
	SD	23.6452	8.8953	14.9517	9.3715	22.4571	7.5480	17.0913	10.1949
	Median	15.5707	16.0192	22.3039	46.1252	16.7361	13.9447	22.2926	41.3245
MI	Mean	0.8885	0.9795	1.2231	2.5453	0.8533	0.8700	1.1704	2.3972
	SD	0.6614	0.5571	0.7276	0.9122	0.6677	0.5209	0.7271	0.8884
	Median	1.1917	1.2528	1.2128	2.8909	1.0294	0.7561	0.9428	2.5234
NMSE	Mean	5.8678	-9.7378	4.3883	-41.1904	12.0661	-8.4432	15.3710	-35.4065
	SD	41.8839	5.9805	47.3965	6.1602	55.2592	5.9695	68.8440	2.4254
	Median	-10.1926	-10.2696	-16.1017	-40.9538	-10.4075	-5.8740	-5.9194	-35.1655
MSG	Mean	20.7649	22.2366	24.8428	54.8803	30.7926	31.2768	34.6683	63.6459
	SD	1.5514	2.8453	5.6492	3.7821	2.5986	2.6873	5.4670	4.6838
	Median	20.4138	21.9056	21.3885	53.9321	29.9387	32.1680	32.1478	64.2189
SNR _{out}	Mean	–	14.5526	–	65.4699	–	23.6387	–	74.2874
	SD	–	2.5991	–	6.6824	–	2.6336	–	7.8450
	Median	–	15.6394	–	61.9911	–	24.2299	–	71.4036

6. Concluding Remarks

In this paper, the main idea is to vary the level of the probe signal so that a fast convergence as well as a good output SNR may be achieved. The simulation results show excellent performance of the proposed method, and it can be promising choice for the practical hearing aids. The proposed method gives plausible performance for speech (Case 1) as well as strongly correlated signals (Case 2). Furthermore, the proposed method shows robust performance against sudden changes in the acoustic feedback path (Case 3). In the future, it would be interesting to investigate the performance of the proposed method for music signals which are considered as a challenge for the present hearing aids. Furthermore, a theoretical analysis of the proposed method is a task for the future work.

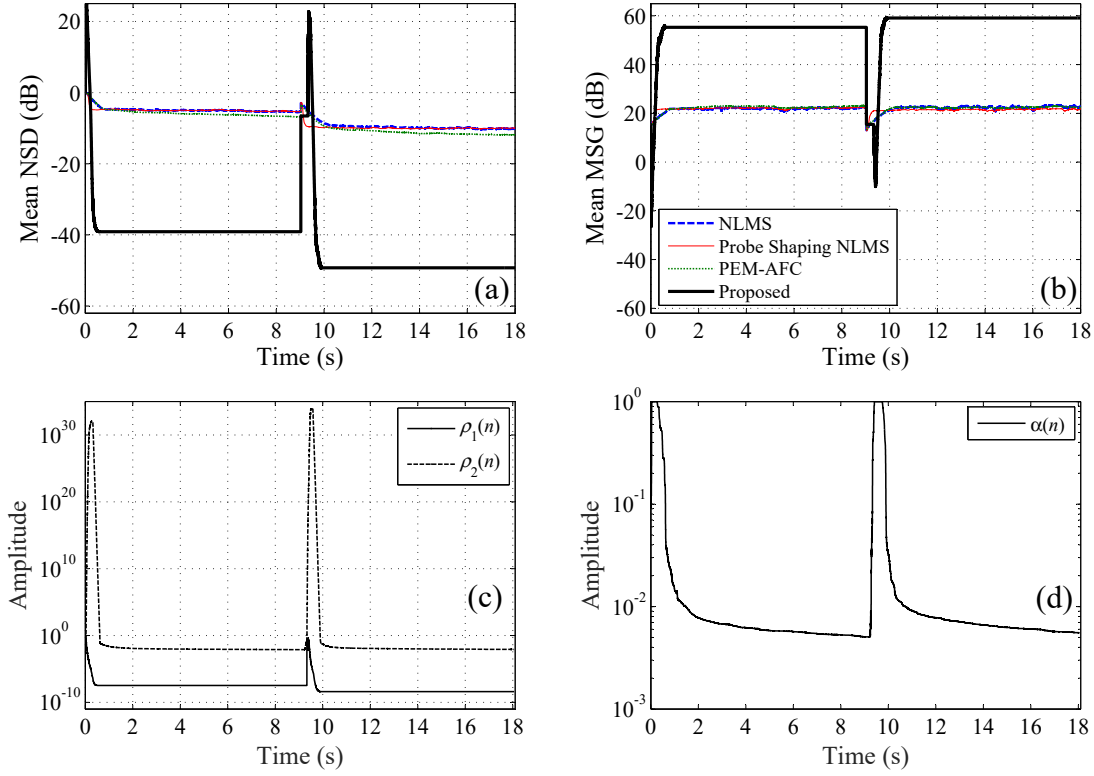


Figure 16: Simulations results in Case 3 (a sudden change in the acoustic feedback path). (a) The curves for mean NSD (in dB), (b) the curves for mean MSG (in dB), (c) variation of parameters $\rho_1(n)$ and $\rho_2(n)$ in the proposed method, and (d) time varying gain $\alpha(n)$ in the proposed method.

Acknowledgments

The authors would like to thank Prof. J. Kates for providing the MATLAB implementations of the HAAQI and HASQI measures. Authors would like to acknowledge anonymous reviewer for providing many critical and insightful comments, which has greatly helped in improving the contents as well as quality of presentation. This research has been funded from the Faculty Development Competitive Research Grants Program of Nazarbayev University under the Grant Number 110119FD4525.

References

- [1] D. K. Bustamante, T. L. Worrall, and M. J. Williamson, "Measurement and adaptive suppression of acoustic feedback in hearing aids," in *Proc. IEEE ICASSP 1989*, pp. 2017–2020.
- [2] J. M. Kates, *Digital Hearing Aids*, Plural Publishing, 2008.
- [3] J. Maxwell and P. Zurek, "Reducing acoustic feedback in hearing aids," *IEEE Trans. Speech Audio Process.*, vol. 4, pp. 304–313, 1995.

- [4] B. W. Edwards, "Signal processing techniques for a DSP hearing aid," in *Proc. IEEE ISCAS 1998*, vol. VI, pp. 586–589.
- [5] A. Kaelin, A. Lindgren, and S. Wyrsh, "A digital frequency domain implementation of a very high gain hearing aid with compensation for recruitment of loudness and acoustic echo cancellation," *Signal Process.*, vol. 64, pp. 71–85, 1998.
- [6] J. M. Kates, "Constrained adaptation for feedback cancellation in hearing aids," *J. Acoust. Soc. Am.*, vol. 106, pp. 1010–1019, 1999.
- [7] S. C. Douglas, "A family of normalized LMS algorithms," *IEEE Signal Process. Lett.*, vol. 1, no. 3, pp. 49–51, 1994.
- [8] M. G. Siqueira, and A. Alwan, "Steady-State analysis of continuous adaptation in acoustic feedback reduction systems for hearing-aids," *IEEE Trans. Speech Audio Process.*, vol. 8, no. 4, pp. 443–453, 2000.
- [9] J. Hellgren, "Analysis of feedback cancellation in hearing aids with filtered-x LMS and the direct method of closed loop identification," *IEEE Trans. Speech Audio Process.*, vol. 10, no. 2, pp. 119–131, 2002.
- [10] M. Guo, S. H. Jensen, and J. Jensen, "Novel acoustic feedback cancellation approaches in hearing aid applications using probe noise and probe noise enhancement," *IEEE Trans. Audio Speech Lang. Process.*, vol. 20, no. 9, pp. 2549–2563, 2012.
- [11] Y. F. Chiang, C. W. Wei, Y. L. Meng, Y. W. Lin, S. J. Jou, and T. S. Chang, "Low complexity formant estimation adaptive feedback cancellation for hearing aids using pitch based processing," *IEEE/ACM Trans. Audio, Speech, Language Process.*, vol. 22, no. 8, pp. 1248–1259, Aug. 2014.
- [12] F. Strasser and H. Puder, "Correlation detection for adaptive feedback cancellation in hearing aids," *IEEE Signal Process. Lett.*, vol. 23, no. 7, pp. 979–983, Jul. 2016.
- [13] P. Estermann and A. Kaelin, "Feedback cancellation in hearing aids: Results from using frequency-domain adaptive filters," in *Proc. IEEE ISCAS 1994*, pp. 257–260.
- [14] A. Pandey and V. J. Mathews, "Low-delay signal processing for digital hearing aids," *IEEE Trans. Audio, Speech, Lang. Process.*, vol. 19, no. 4, pp. 699–710, May 2011.
- [15] H. Sakai and H. Fukuzono, "Analysis of adaptive filters in feedback cancellation for sinusoidal signals," in *Proc. APSIPA-ASC 2009*, pp. 430–433.
- [16] A. Spriet, I. Proudler, M. Moonen, and J. Wouters, "Adaptive feedback cancellation in hearing aids with linear prediction of the desired signal," *IEEE Trans. Signal Process.*, vol. 53, no. 10, pp. 3749–3763, 2005.
- [17] G. Rombouts, T. Van Waterschoot, and M. Moonen, "Robust and efficient implementation of the PEM AFROW algorithm for acoustic feedback cancellation," *J. Audio Eng. Soc.*, vol. 55, no. 11, pp. 955–966, 2007.
- [18] A. Spriet, S. Doclo, M. Moonen, and J. Wouters, "Feedback control in hearing aids," in *Springer Handbook of Speech Processing*. J. Benesty, M. Sondhi, and Y. Huang, Eds. Berlin, Germany: Springer, 2008, pp. 979–1000.
- [19] R. Vican-Bueno, A. Martínez-Leira, R. Gil-Pita, and M. Rosa-Zurera, "Modified LMS-based feedback-reduction subsystems in digital hearing aids based on WOLA filter bank," *IEEE Trans. Instrum. Meas.*, vol. 58, no. 9, pp. 3177–3190, 2009.
- [20] C. R. C. Nakagawa, S. Nordholm, and W. Y. Yan, "Dual microphone solution for acoustic feedback cancellation for assistive learning," in *Proc. IEEE ICASSP 2012*, pp. 149–152.
- [21] C. R. C. Nakagawa, S. Nordholm, and W. Y. Yan, "Analysis of two microphone method for feedback cancellation," *IEEE Signal Process. Lett.*, vol. 22, no. 1, pp. 35–39, 2015.
- [22] J. E. Greenberg, P. M. Zurek, and M. Brantley, "Evaluation of feedback-reduction algorithms for hearing aids," *J. Acoust. Soc. Am.*, vol. 108, no. 5, pp. 2366–2376, 2000.
- [23] C. R. C. Nakagawa, S. Nordholm, and W. Y. Yan, "Feedback cancellation with probe shaping compensation," *IEEE Signal Process. Lett.*, vol. 21, no. 3, pp. 365–369, 2014.
- [24] A. Mader, H. Puder, and G. U. Schmidt, "Step-size control for acoustic echo cancellation filters—An overview," *Signal Process.*, vol. 4, pp. 1697–1719, 2000.
- [25] M. T. Akhtar, and A. Nishihara, "Automatic tuning of probe noise for continuous acoustic feedback cancellation in hearing

- aids,” in *Proc. EUSIPCO 2016*, August 29 - September 2, 2016, Budapest Hungary, pp. 888–892.
- [26] H. Schepker, “Robust feedback suppression algorithms for single- and multimicrophone hearing aids,” *PhD Thesis*, University of Oldenburg, Germany, Dec. 2017.
 - [27] J. Hellgren, and U. Forssell, “Bias of feedback cancellation algorithms in hearing aids based on direct closed loop identification,” *IEEE Trans. Speech Audio Process.*, vol. 9, no. 7, pp. 906–913, Nov. 2001.
 - [28] G. Rombouts, T. van Waterschoot, K. Struyve, and M. Moonen, “Acoustic feedback cancellation for long acoustic paths using a nonstationary source model,” *IEEE Trans. Signal Process.*, vol. 54, no. 9, pp. 3426–3434, Sep. 2006.
 - [29] A. Spriet, G. Rombouts, M. Moonen, and J. Wouters, “Adaptive feedback cancellation in hearing aids,” *J. Franklin Inst.*, vol. 343, no. 6, pp. 545–573, Aug. 2006.
 - [30] J. M. Gil-Cacho, T. van Waterschoot, M. Moonen, and S. H. Jensen, “Wiener variable step size and gradient spectral variance smoothing for double-talk-robust acoustic echo cancellation and acoustic feedback cancellation,” *Signal Process.*, vol. 104, pp. 1–14, Nov. 2014.
 - [31] G. Bernardi, T. van Waterschoot, J. Wouters, and M. Moonen, “Adaptive feedback cancellation using a partitioned-block frequency-domain Kalman filter approach with PEM-based signal prewhitening,” *IEEE/ACM Trans. Audio Speech Lang. Process.*, vol. 25, no. 9, pp. 1480–1494, Sep. 2017.
 - [32] M. Hayes, *Statistical Digital Signal Processing and Modeling*, Wiley, 1996.
 - [33] T. Sankowsky-Rothe, and M. Blau, “Static and dynamic measurements of the acoustic feedback path of hearing aids on human subjects,” *Proc. Mtgs. Acoust.*, vol. 30, no. 050008, 2017. (doi: 10.1121/2.0000618)
 - [34] R. Heurig, and J. Chalupper, “Acceptable processing delay in digital hearing aids,” *Hearing Review*, vol. 17, no. 1, pp. 28–31, 2010.
 - [35] J. M. Kates and K. H. Arehart, “The hearing-aid speech quality index (HASQI),” *J. Audio Eng. Soc.*, vol. 58, no. 5, pp. 363–381, 2010.
 - [36] J. M. Kates and K. H. Arehart, “The hearing-aid speech quality index (HASQI) version 2,” *J. Audio Eng. Soc.*, vol. 62, no. 3, pp. 99–117, 2014.
 - [37] J. M. Kates and K. H. Arehart, “The hearing-aid audio quality index (HAAQI),” *IEEE/ACM Trans. Audio Speech Lang. Process.*, vol. 24, no. 2, pp. 354–365, Feb. 2016.
 - [38] P. C. Loizou, *Speech Enhancement Theory and Practice*, CRC Press, 2007.
 - [39] A. J. Manders, D. M. Simpson, and S. L. Bell, “Objective prediction of the sound quality of music processed by an adaptive feedback canceller,” *IEEE Trans. Audio Speech Lang. Process.*, vol. 20, no. 6, pp. 1734–1745, 2012.
 - [40] A. Olofsson and M. Hansen, “Objectively measured and subjectively perceived distortion in nonlinear systems,” *J. Acoust. Soc. Amer.*, vol. 120, no. 6, pp. 3759–3769, 2006.
 - [41] E. Vincent, R. Gribonval, and C. Févotte, “Performance measurement in blind audio source separation,” *IEEE Trans. Audio, Speech Lang. Process.*, vol. 14, no. 4, pp. 1462–1469, 2006.
 - [42] A. Hyvarinen and E. Oja, “Independent component analysis: algorithms and applications,” *Neural Networks*, vol. 13, no. 4–5, pp. 411–430, 2000.
 - [43] G. Bernardi, T. van Waterschoot, J. Wouters, and Marc Moonen, “Subjective and objective sound-quality evaluation of adaptive feedback cancellation algorithms,” *IEEE/ACM Trans. Audio Speech Lang. Process.*, vol. 26, no. 5, pp. 1010–1024, May 2018.
 - [44] S. Banerjee, K. Recker, and A. Paumen, “A tale of two feedback cancellers,” *Hearing Review*, Jul., 2006. (Available online at http://www.hearingreview.com/issues/articles/2006-07_29.asp)
 - [45] M. T. Akhtar, and Akinori Nishihara, “Howling and Entrainment in Hearing Aids: A Review,” *Engineering Journal*, vol. 20, no. 5, pp. 5–13, 2016. (<http://engj.org/index.php/ej/article/view/1321/497>)