

RANK-1 APPROXIMATION BASED MULTICHANNEL WIENER FILTERING ALGORITHMS FOR NOISE REDUCTION IN COCHLEAR IMPLANTS

Romain Serizel¹, Marc Moonen¹, Bas Van Dijk² and Jan Wouters³

¹KU Leuven, ESAT-SCD and iMinds Future Health Department,
Kasteelpark Arenberg 10, B-3001 Leuven, Belgium

²Cochlear CTCE, Schalinhoevedreef 20, Building i-B, B-2800 Mechelen, Belgium

³KU Leuven, ExpORL, O.& N2, Herestraat 49/721, B-3000 Leuven, Belgium

ABSTRACT

This paper presents multichannel Wiener filtering-based algorithms for noise reduction in cochlear implants. In a single speech scenario, the autocorrelation matrix of the speech signal can be approximated by a rank-1 matrix. It is then possible to derive noise reduction filters that deliver improved signal-to-noise ratio performance. The link between these different filters is investigated here and an eigenvalue decomposition based algorithm is demonstrated to be more stable at low input signal-to-noise ratio compared to previous algorithms.

Index Terms— Speech enhancement, multichannel Wiener filter, cochlear implant, rank-1, eigenvalues decomposition

1. INTRODUCTION

A major challenge in cochlear implant (CI) is to improve the speech understanding in noise [1] and so having an efficient front-end noise reduction (NR) is important. Therefore, during the past years, several NR algorithms have been developed and tested with CI recipients [2, 3, 4].

In general, CI users need a 10dB to 25dB higher signal-to-noise ratio (SNR) than normal hearing listeners to achieve similar speech understanding performance [5]. This could motivate the use of more aggressive NR strategies. Speech distortion weighted multichannel Wiener filters (SDW-MWF) have been developed to allow to tune multichannel Wiener filter (MWF)-based NR and perform a more aggressive NR by allowing more speech distortion (SD) [6, 7, 8]. In the case of a single speech source the SDW-MWF performance can be improved if the filters are reformulated based on the assumption that the speech autocorrelation matrix is rank-1, leading, *e.g.*, to the so-called spatially-preprocessed MWF (SP-MWF) [9, 10].

All these NR algorithms rely on the estimation of the speech autocorrelation matrix. At low SNR, the speech autocorrelation matrix can be wrongly estimated and become non positive semi-definite. The SDW-MWF and the SP-MWF can then behave unpredictably. A solution to this problem is then to select a rank-1 approximation based on an eigenvalue decomposition (EVD) of the speech autocorrelation matrix. This paper also presents a performance comparison

between the original SDW-MWF and the EVD-based NR applied on both bilateral and binaural set-ups [11, 12, 13, 14].

The signal model and the SDW-MWF-based NR are described in Section 3. Section 4 describes the so-called first column decomposition and how this provides an interpretation of the SDW-MWF and the SP-MWF. The EVD-based NR is introduced in Section 5. The performance of the original SDW-MWF and the EVD-based NR are compared in Section 6. Finally, Section 7 presents a summary of the paper.

2. RELATION TO PRIOR WORK

The work presented in this paper focuses on the analysis of the difference between the SDW-MWF [6, 7, 8] and the SP-MWF [9, 10] when the rank of the estimated autocorrelation matrix of the speech signal is higher than one. A new rank-1 approximation is also introduced. While the (implicit or explicit) rank-1 approximation in the previous work was based on a so-called first column decomposition, the new (explicit) rank-1 approximation presented here is based on the EVD.

3. BACKGROUND AND PROBLEM STATEMENT

3.1. Signal model

Let M be the number of microphones (channels). The frequency-domain signal $X_m(\omega)$ for microphone m has a desired speech part $X_{m,s}(\omega)$ and an additive noise part $X_{m,n}(\omega)$, *i.e.*:

$$X_m(\omega) = X_{m,s}(\omega) + X_{m,n}(\omega) \quad m \in \{1 \dots M\} \quad (1)$$

where $\omega = 2\pi f$ is the frequency-domain variable. For conciseness, ω will be omitted in all subsequent equations.

Signal model (1) holds for so-called “speech plus noise periods”. There are also “noise only periods” (*i.e.*, speech pauses), during which only a noise component is observed.

In practice, in order to distinguish between “speech plus noise periods” and “noise only periods” it is necessary to use a voice activity detector (VAD). The performance of the VAD can affect the performance of the noise reduction. For the time being, a perfect VAD is assumed.

The compound vector gathering all microphone signals is:

$$\mathbf{X} = [X_1 \dots X_M]^T \quad (2)$$

An optimal (Wiener) filter $\mathbf{W} = [W_1 \dots W_M]^T$ will be designed and applied to the signals, which minimizes a Mean Squared Error

This research work was carried out at the ESAT Laboratory of KU Leuven, in the frame of KU Leuven Research Council CoE EF/05/006 Optimization in Engineering (OPTeC), IWT Project ‘Signal processing and automatic fitting for next generation cochlear implants’, PFV/10/002 (OPTeC), Concerted Research Action GOA-MaNet, the Belgian Programme on Interuniversity Attraction Poles initiated by the Belgian Federal Science Policy Office IUAP P6/04 (DYSCO, ‘Dynamical systems, control and optimization’, 2007-2011), Research Project FWO nr. G.0600.08 (‘Signal processing and network design for wireless acoustic sensor networks’), EC-FP6 project SIGNAL: ‘Core Signal Processing Training Program’. The scientific responsibility is assumed by its authors.

(MSE) criterion:

$$J_{\text{MSE}} = \mathbb{E}\{|E|^2\} \quad (3)$$

Where $\mathbb{E}$ is the expectation operator and E is an error signal to be defined next, depending on the scheme applied. The filter output signal Z is defined as:

$$Z = \mathbf{W}^H \mathbf{X} \quad (4)$$

where H denotes the Hermitian transpose.

The desired speech signal is arbitrarily chosen to be the (unknown) speech component of the first microphone signal ($m = 1$). This can be written as:

$$D_{\text{NR}} = \mathbf{e}_1^H \mathbf{X}_s \quad (5)$$

where $\mathbf{e}_1$ is an all-zero vector expect for a one in the first position.

The autocorrelation matrices of the microphone signals in “speech plus noise periods”, and of the speech component and the noise component of the microphone signals are given by:

$$\mathbf{R}_x = \mathbb{E}\{\mathbf{X}\mathbf{X}^H\} \quad (6)$$

$$\mathbf{R}_s = \mathbb{E}\{\mathbf{X}_s \mathbf{X}_s^H\} \quad (7)$$

$$\mathbf{R}_n = \mathbb{E}\{\mathbf{X}_n \mathbf{X}_n^H\} \quad (8)$$

$\mathbf{R}_n$ can be estimated during “noise only periods” and $\mathbf{R}_x$ can be estimated during “speech plus noise periods”. If the speech and noise signals are assumed to be uncorrelated and if the noise signal is stationary, $\mathbf{R}_s$ can be estimated by using:

$$\mathbf{R}_s = \mathbf{R}_x - \mathbf{R}_n \quad (9)$$

3.2. MWF-based Noise Reduction

The MWF aims to minimize the squared distance between the filtered microphone signal and the desired speech signal. The corresponding MSE criterion is:

$$J_{\text{MSE}} = \mathbb{E}\{|E_{\text{MWF}}|^2\} \quad (10)$$

$$E_{\text{MWF}} = \mathbf{W}^H \mathbf{X} - \mathbf{e}_1^H \mathbf{X}^s \quad (11)$$

The MWF solution is given as:

$$\mathbf{W}_{\text{MWF}} = (\mathbf{R}_s + \mathbf{R}_n)^{-1} \mathbf{R}_s \mathbf{e}_1 \quad (12)$$

The SDW-MWF has been proposed to provide an explicit trade-off between the SD and the NR [6, 7, 8]:

$$J_{\text{MSE}} = \mathbb{E}\{|\mathbf{W}^H \mathbf{X}^s - \mathbf{e}_1^H \mathbf{X}^s|^2\} + \mu \mathbb{E}\{|\mathbf{W}^H \mathbf{X}^n|^2\} \quad (13)$$

The SDW-MWF solution is then given as:

$$\mathbf{W}_{\text{SDW-MWF}} = (\mathbf{R}_s + \mu \mathbf{R}_n)^{-1} \mathbf{R}_s \mathbf{e}_1 \quad (14)$$

In a single speech source scenario, the autocorrelation matrix of the speech component of the microphone signals $\mathbf{R}_s$ is often assumed to be a rank-1 matrix, which can then be rewritten as:

$$\mathbf{R}_s = P^s \mathbf{A} \mathbf{A}^H \quad (15)$$

where P^s is the power of the speech source signal and $\mathbf{A}$ is the M -dimensional steering vector, containing the acoustic transfer functions from the speech source position to the hearing aid microphones (including room acoustics, microphone characteristics, and head shadow effect).

Based on this rank-1 assumption it is possible to derive the SP-MWF [9, 10]:

$$\mathbf{W}_{\text{SP-MWF}} = \mathbf{R}_n^{-1} \mathbf{R}_s \mathbf{e}_1 \frac{\mathbf{e}_1^H \mathbf{R}_s \mathbf{e}_1}{\mu \mathbf{e}_1^H \mathbf{R}_s \mathbf{e}_1 + \text{Tr}\{\mathbf{R}_n^{-1} \mathbf{R}_s \mathbf{e}_1 \mathbf{e}_1^H \mathbf{R}_s\}} \quad (16)$$

The filters (14) and (16) are fully equivalent if $\mathbf{R}_s$ is rank-1. In practice, however $\text{rank}(\mathbf{R}_s) > 1$ even for a single speech source scenario and then (14) and (16) are different filters.

4. FIRST COLUMN DECOMPOSITION

When $\text{rank}(\mathbf{R}_s) > 1$ the speech autocorrelation matrix $\mathbf{R}_s$ can still be decomposed as:

$$\mathbf{R}_s = \mathbf{R}_{s_{r1}} + \mathbf{R}_Z \quad (17)$$

where $\mathbf{R}_{s_{r1}}$ is a rank-1 matrix and $\mathbf{R}_Z$ is a “remainder” matrix.

The most obvious choice for $\mathbf{R}_{s_{r1}}$ is then a rank-1 extension of its first column and row, *i.e.*:

$$\mathbf{R}_s = \underbrace{\mathbf{d} \mathbf{d}^H}_{\mathbf{R}_{s_{r1}}} + \mathbf{R}_Z \quad (18)$$

where

$$\sigma_{i,j} = [\mathbf{R}_s]_{i,j} \quad (19)$$

$$\mathbf{d} = [1 \frac{\sigma_{2,1}}{\sigma_{1,1}} \dots \frac{\sigma_{1,N}}{\sigma_{1,1}}]^T \quad (20)$$

$$\mathbf{R}_Z = \begin{bmatrix} 0 & 0 & \dots & 0 \\ 0 & x & \dots & x \\ \vdots & \vdots & & \vdots \\ 0 & x & \dots & x \end{bmatrix} \quad (21)$$

and $\sigma_{1,1}$ is the speech power in microphone 1.

It is noted that:

$$\mathbf{R}_s \mathbf{e}_1 = \mathbf{R}_{s_{r1}} \mathbf{e}_1 + \underbrace{\mathbf{R}_Z \mathbf{e}_1}_{=0} \quad (22)$$

which means that the (rightmost) “desired signal part” $\mathbf{R}_s \mathbf{e}_1$ in (14) and (16) can (obviously) be replaced by the “rank-1 desired signal part” $\mathbf{R}_{s_{r1}} \mathbf{e}_1$. The difference between the two approaches (14) and (16) then effectively depends on how $\mathbf{R}_Z$ is treated.

4.1. SDW-MWF

Plugging (18) into the SDW-MWF formula (14) leads to:

$$\mathbf{W}_{\text{SDW-MWF}} = (\mathbf{R}_{s_{r1}} + \mu(\mathbf{R}_n + \frac{1}{\mu} \mathbf{R}_Z))^{-1} \mathbf{R}_{s_{r1}} \mathbf{e}_1 \quad (23)$$

This means that in the SDW-MWF $\mathbf{R}_s$ is replaced by $\mathbf{R}_{s_{r1}}$ and the remainder matrix $\mathbf{R}_Z$ is effectively treated as noise (up to a scaling with $\frac{1}{\mu}$).

4.2. SP-MWF

Plugging (18) into the SP-MWF formula (16) leads to:

$$\mathbf{W}_{\text{SP-MWF}} = (\mathbf{R}_{s_{r1}} + \mu \mathbf{R}_n)^{-1} \mathbf{R}_{s_{r1}} \mathbf{e}_1 \quad (24)$$

This means that the SP-MWF corresponds to the SDW-MWF (14) where $\mathbf{R}_s$ is replaced by $\mathbf{R}_{s_{r1}}$ and the remainder matrix $\mathbf{R}_Z$ is simply ignored. If $\mathbf{R}_Z = 0$ (rank-1 case) formulas (23) and (24) are again seen to be the same.

4.3. Speech autocorrelation matrix estimation

At low SNR and if the noise is non-stationary it is observed that:

$$\mathbf{R}_x \approx \mathbf{R}_n \quad (25)$$

and then the estimated speech autocorrelation matrix $\mathbf{R}_s = \mathbf{R}_x - \mathbf{R}_n$ can lose its positive definiteness, which has been observed to lead to filter instabilities. The first column decomposition-based filters suffer from the same estimation problem where then the estimated speech power in microphone 1 ($\sigma_{1,1}$) could become negative, *i.e.*, $\mathbf{R}_{s,r1}$ is non-positive definite and so that the desired signal is ill-defined.

5. EVD-BASED FILTERS

An alternative to the first column decomposition based rank-1 approximation would be to consider a rank-1 approximation based on an EVD of $\mathbf{R}_s$:

$$\mathbf{R}_s = \underbrace{\mathbf{d}_{\max} \mathbf{d}_{\max}^H \lambda_{\max}}_{\mathbf{R}_{s,r1}} + \mathbf{R}_Z \quad (26)$$

where $\lambda_{\max}$ is $\mathbf{R}_s$'s largest eigenvalue and $\mathbf{d}_{\max}$ is the corresponding normalized eigenvector and $\mathbf{R}_Z$ is again a remainder matrix. In this case, the rank-1 part $\mathbf{R}_{s,r1}$ is positive definite if the dominant eigenvalue of $\mathbf{R}_s$ is positive (which is more likely than the first diagonal element $\sigma_{1,1}$ of the matrix $\mathbf{R}_s$ being positive).

All the formulas introduced above can be modified based on this decomposition.

It is noted that now:

$$\mathbf{R}_s \mathbf{f}_1 = \mathbf{R}_{s,r1} \mathbf{f}_1 + \underbrace{\mathbf{R}_Z \mathbf{f}_1}_{=0} \quad (27)$$

$$\mathbf{R}_{s,r1} \mathbf{f}_1 = \mathbf{R}_{s,r1} \mathbf{e}_1 \quad (28)$$

to be compared to (22) where

$$\mathbf{f}_1 = \mathbf{d}_{\max} \mathbf{d}_{\max}^H(1)^* \quad (29)$$

An analysis similar to the analysis of the first column decomposition in Section 4 can then be done where $\mathbf{R}_s$ is replaced by $\mathbf{R}_{s,r1}$ and the remainder matrix $\mathbf{R}_Z$ is either treated as noise or ignored. Equivalently, one can start from a modified SDW-MWF criterion where the (arbitrary) $\mathbf{e}_1$ is replaced by $\mathbf{f}_1$:

$$J_{\text{MSE}} = \mathbb{E}\{|\mathbf{W}^H \mathbf{X}^s - \mathbf{f}_1^H \mathbf{X}^s|^2\} + \mu \mathbb{E}\{|\mathbf{W}^H \mathbf{X}^n|^2\} \quad (30)$$

5.1. SDW-MWF_{EVD}

Plugging (27) into the SDW-MWF formula corresponding to (30) leads to:

$$\mathbf{W}_{\text{SDW-MWF}} = (\mathbf{R}_{s,r1} + \mu(\mathbf{R}_n + \frac{1}{\mu} \mathbf{R}_Z))^{-1} \underbrace{\mathbf{R}_{s,r1} \mathbf{f}_1}_{\mathbf{R}_{s,r1} \mathbf{e}_1} \quad (31)$$

This means that in the SDW-MWF $\mathbf{R}_s$ is replaced by the EVD-based $\mathbf{R}_{s,r1}$ and the remainder matrix $\mathbf{R}_Z$ is effectively treated as noise (up to a scaling with $\frac{1}{\mu}$).

5.2. SP-MWF_{EVD}

Plugging (27) into the SP-MWF formula corresponding to (30) leads to:

$$\mathbf{W}_{\text{SP-MWF}} = (\mathbf{R}_{s,r1} + \mu \mathbf{R}_n)^{-1} \underbrace{\mathbf{R}_{s,r1} \mathbf{f}_1}_{\mathbf{R}_{s,r1} \mathbf{e}_1} \quad (32)$$

This means that in the SP-MWF $\mathbf{R}_s$ is replaced by the EVD-based $\mathbf{R}_{s,r1}$ and the remainder matrix $\mathbf{R}_Z$ is simply ignored.

6. EXPERIMENTAL RESULTS

6.1. Experimental setup

The simulations were run on acoustic path measurements obtained in a reverberant room ($\text{RT}_{60} = 0.61\text{s}$ [15, 16]) with a CORTEX MK2 manikin head and torso equipped with two Cochlear SP15 behind-the-ear devices. Each device has two omnidirectional microphones. The sound sources (FOSTEX 6301B loudspeakers) were positioned at 1 meter from the center of the head.

The speech signal was composed of five consecutive sentences from the English Hearing-In-Noise Test (HINT) database [17] concatenated with five seconds silence periods. The noise was the multitalker babble from Auditec [18]. The speech source was located at 0° and the noise source at 45° (S0N45). All the signals were sampled at 20480Hz. The filter lengths and DFT size were set to $N = 128$ and the frame overlap was set to half of the DFT size ($L = 64$). When mentioned, the so-called input SNR is the SNR at the center of the head (excluding the head shadow effects).

6.2. Performance measures

The speech intelligibility-weighted SNR (SIW-SNR) [19] is used here to compute the *SIW-SNR improvement* which is defined as

$$\Delta \text{SNR}_{\text{intellig}} = \sum_i I_i (\text{SNR}_{i,\text{out}} - \text{SNR}_{i,\text{ref}}) \quad (33)$$

where I_i is the band importance function defined in [20] and $\text{SNR}_{i,\text{out}}$ and $\text{SNR}_{i,\text{ref}}$ represent the output SNR (at the considered ear) of one of the NR schemes and the SNR of the signal in the reference microphone (at the considered ear) of the i th band, respectively.

6.3. S0N45

Signals with input SNR varying from -15dB to 5dB are presented to the left and right hearing aid devices. The microphone signals are then filtered by several NR algorithms and the performance is compared.

The SIW-SNR improvement at the left ear for bilateral NR filters is presented in Figure 1. The EVD-based rank-1 filters improve the SIW-SNR by about 2dB compared to the SDW-MWF. At low SNR, the behaviour of the SP-MWF is unpredictable, this is caused by the sensitivity of the SP-MWF to $\mathbf{R}_s$'s that are not positive definite.

The next results demonstrate how the EVD-based NR can benefit from binaural setups. Three setups are considered: the bilateral setup (2 microphones), a so-called binaural "front" setup where only the signal from the front microphone of the contra-lateral ear is shared (2+1 microphones) and a setup where both microphone signals from the contra-lateral ear are shared (2+2 microphones, this approach is referred to as "binaural" here).

Figure 2 presents percentage of estimated $\mathbf{R}_s$'s that are not positive definite, at the left ear, as a function of the input SNR for bilateral, front and binaural SDW-MWF and SDW-MWF_{EVD}. In the

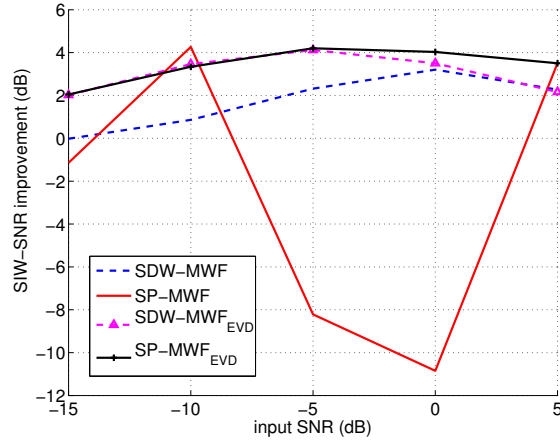


Fig. 1. SIW-SNR performance (bilateral NR filters)

SDW-MWF based NR the positive definiteness of $\mathbf{R}_{s,r-1}$ only depends on the first diagonal element of the speech autocorrelation matrix $\mathbf{R}_s(1, 1)$, therefore, bilateral, front and binaural approaches return the same percentage of estimated $\mathbf{R}_s$'s that are not positive definite which can be as high as 65% at low SNR. In the SDW-MWF_{EVD} on the other hand, the positive definiteness of $\mathbf{R}_{s,r-1}$ depends on $\lambda_{\max}$ and each additional channel can help increasing this quantity. Therefore, whereas the bilateral SDW-MWF_{EVD} can help to decrease the percentage of estimated $\mathbf{R}_s$'s that are not positive definite from about 65% to 60% at low SNR, the front and the binaural SDW-MWF_{EVD} allow to decrease this percentage to around 40%.

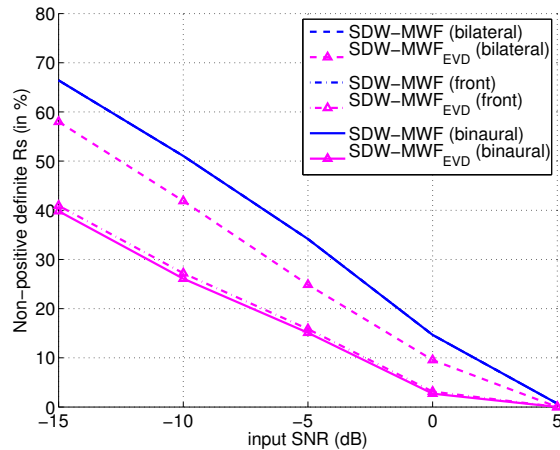


Fig. 2. Percentage of estimated $\mathbf{R}_s$'s that are not positive definite (SDW-MWF and SDW-MWF_{EVD})

Figures 3 presents the SIW-SNR improvement at the left ear. The front and the binaural SDW-MWF allow to improve the SIW-SNR from 2dB to 6dB depending on the input SNR. The bilateral SDW-MWF_{EVD} provides an SIW-SNR improvement from 2dB to 4dB for any input SNR. This is 2dB better than the SIW-SNR improvement of the bilateral SDW-MWF. The front and the binaural SDW-MWF_{EVD} provide an SIW-SNR improvement between 6dB

and 12dB depending on the input SNR. This is 4dB to 6dB better than with the respective SDW-MWF approaches.

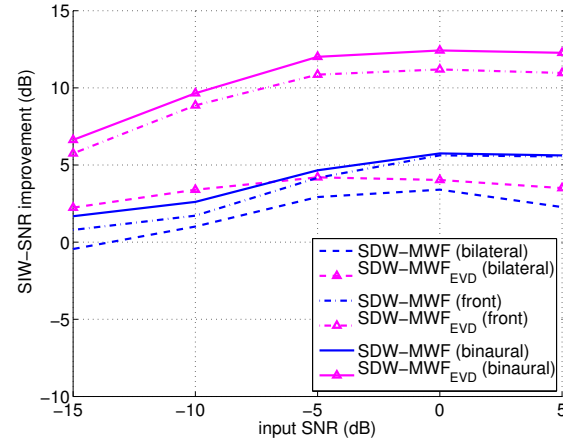


Fig. 3. SIW-SNR performance (SDW-MWF and SDW-MWF_{EVD})

7. CONCLUSION

In this paper the difference between the SDW-MWF and the SP-MWF (which are equivalent when the autocorrelation matrix of the speech signal is a rank-1 matrix) is analysed when the rank of the autocorrelation matrix of the speech signal is higher than one. In this case, it is possible to decompose the autocorrelation matrix of the speech signal into the sum of a rank-1 matrix and a remainder matrix. The SDW-MWF and the SP-MWF then differs in the way this remainder matrix is treated.

At low input SNR, due to noise non-stationarity, the speech autocorrelation matrix may become non-positive definite. An EVD-based rank-1 approach to SDW-MWF and to SP-MWF has been introduced. It is then again possible to decompose the autocorrelation matrix of the speech signal into the sum of a rank-1 matrix and a remainder matrix and the difference between the EVD-based SDW-MWF and SP-MWF depends in the way the remainder matrix is treated. It is demonstrated that the SDW-MWF_{EVD} allows an improved SIW-NSR performance.

8. REFERENCES

- [1] Y. Hu and P.C. Loizou, "A new sound coding strategy for suppressing noise in cochlear implants," *The Journal of the Acoustical Society of America*, vol. 124, no. 1, pp. 498, 2008.
- [2] RJM Van Hoesel and GM Clark, "Evaluation of a portable two-microphone adaptive beamforming speech processor with cochlear implant patients," *The Journal of the Acoustical Society of America*, vol. 97, pp. 2498, 1995.
- [3] M.R. Weiss et al., "Effects of noise and noise reduction processing on the operation of the nucleus-22 cochlear implant processor," *Journal of rehabilitation research and development*, vol. 30, pp. 117-117, 1993.
- [4] A. Spriet, L. Van Deun, K. Eftaxiadis, J. Laneau, M. Moonen, B. van Dijk, A. Van Wieringen, and J. Wouters, "Speech understanding in background noise with the two-microphone

- adaptive beamformer beam (tm) in the nucleus freedom (tm) cochlear implant system,” *Ear and Hearing*, vol. 28, no. 1, pp. 62–72, 2007.
- [5] J. Wouters and J.V. Berghe, “Speech recognition in noise for cochlear implantees with a two-microphone monaural adaptive noise reduction system,” *Ear and Hearing*, vol. 22, no. 5, pp. 420–430, 2001.
 - [6] S. Doclo, A. Spriet, J. Wouters, and M. Moonen, “Frequency-domain criterion for the speech distortion weighted multichannel wiener filter for robust noise reduction,” *Elsevier Speech Communication*, vol. 49, no. 7-8, pp. 636–656, 2007.
 - [7] Y. Ephraim and H. L. Van Trees, “A signal subspace approach for speech enhancement,” *IEEE Transactions on Speech and Audio Processing*, vol. 3, no. 4, pp. 251–266, 1995.
 - [8] A. Spriet, M. Moonen, and J. Wouters, “Spatially pre-processed speech distortion weighted multi-channel wiener filtering for noise reduction,” *EURASIP Signal Processing*, vol. 84, no. 12, pp. 2367–2387, 2004.
 - [9] B. Cornelis, M. Moonen, and J. Wouters, “Performance analysis of multichannel wiener filter-based noise reduction in hearing aids under second order statistics estimation errors,” *Audio, Speech, and Language Processing, IEEE Transactions on*, vol. 19, no. 5, pp. 1368–1381, 2011.
 - [10] J. Benesty, J. Chen, and Y. Huang, “Noncausal (frequency-domain) optimal filters,” in *Microphone Array Signal Processing*, vol. 1 of *Springer Topics in Signal Processing*, pp. 115–137. Springer Berlin Heidelberg, 2008.
 - [11] Y. Hamacher, “Comparison of advanced monaural and binaural noise reduction algorithms for hearing aids,” in *IEEE International Conference on Acoustics, Speech and Signal Processing, (ICASSP’02)*. IEEE, 2002, vol. 4.
 - [12] T. J. Klasen, T. Van den Bogaert, M. Moonen, and J. Wouters, “Binaural noise reduction algorithms for hearing aids that preserve interaural time delay cues,” *IEEE Transactions on Signal Processing*, vol. 55, no. 4, pp. 1579–1585, 2007.
 - [13] A. Markides, *Binaural hearing aids*, Academic Press London, 1977.
 - [14] T. Van den Bogaert, *Preserving binaural cues in noise reduction algorithms for hearing aids.*, Ph.D. thesis, Katholieke Universiteit Leuven, Leuven, Belgium, Faculty of Engineering, 2008.
 - [15] T. Van den Bogaert, S. Doclo, M. Moonen, and J. Wouters, “The effect of multimicrophone noise reduction systems on sound source localization by users of binaural hearing aids,” *The Journal of the Acoustical Society of America*, vol. 124, no. 1, pp. 484–497, 2008.
 - [16] T. Van den Bogaert, S. Doclo, J. Wouters, and M. Moonen, “Speech enhancement with multichannel wiener filter techniques in multimicrophone binaural hearing aids,” *The Journal of the Acoustical Society of America*, vol. 125, pp. 360, 2009.
 - [17] M. Nilsson, S. D. Soli, and A. Sullivan, “Development of the Hearing in Noise Test for the measurement of speech reception thresholds in quiet and in noise,” *The Journal of the Acoustical Society of America*, vol. 95, no. 2, pp. 1085–1099, Feb. 1994.
 - [18] Auditec, “Auditory Tests (Revised), Compact Disc, Auditec, St. Louis,” 1997.
 - [19] J. E. Greenberg, P. M. Peterson, and P. M. Zurek, “Intelligibility-weighted measures of speech-to-interference ratio and speech system performance,” *The Journal of the Acoustical Society of America*, vol. 94, no. 5, pp. 3009–3010, Nov. 1993.
 - [20] ANSI-SII (1997), “ANSI S3.5-1997 American National Standard Methods for calculation of the speech intelligibility index,” *Acoustical Society of America*, June 1997.