

Reasons why Current Speech-Enhancement Algorithms do not Improve Speech Intelligibility and Suggested Solutions

Philipos C. Loizou, *Senior Member, IEEE*, and Gibak Kim

Abstract—Existing speech enhancement algorithms can improve speech quality but not speech intelligibility, and the reasons for that are unclear. In the present paper, we present a theoretical framework that can be used to analyze potential factors that can influence the intelligibility of processed speech. More specifically, this framework focuses on the fine-grain analysis of the distortions introduced by speech enhancement algorithms. It is hypothesized that if these distortions are properly controlled, then large gains in intelligibility can be achieved. To test this hypothesis, intelligibility tests are conducted with human listeners in which we present processed speech with controlled speech distortions. The aim of these tests is to assess the perceptual effect of the various distortions that can be introduced by speech enhancement algorithms on speech intelligibility. Results with three different enhancement algorithms indicated that certain distortions are more detrimental to speech intelligibility degradation than others. When these distortions were properly controlled, however, large gains in intelligibility were obtained by human listeners, even by spectral-subtractive algorithms which are known to degrade speech quality and intelligibility.

Index Terms—Ideal binary mask, speech distortions, speech enhancement, speech intelligibility improvement.

I. INTRODUCTION

MUCH progress has been made in the development of speech enhancement algorithms capable of improving speech quality [1], [2]. In stark contrast, little progress has been made in designing algorithms that can improve speech intelligibility. The first intelligibility study done by Lim [3] in the late 1970s found no intelligibility improvement with the spectral subtraction algorithm for speech corrupted in white noise at -5 to 5 dB signal-to-noise ratio (SNR). In the intelligibility study by Hu and Loizou [4], conducted 30 years later, none of the eight different algorithms examined were found to improve speech intelligibility relative to unprocessed (corrupted) speech. Noise reduction algorithms implemented in wearable hearing aids revealed no significant intelligibility benefit, but improved ease of listening and listening comfort [5] for hearing-

impaired listeners. In brief, the ultimate goal of devising an algorithm that would improve speech intelligibility for normal-hearing or hearing-impaired listeners has been elusive for nearly three decades.

Little is known as to why speech enhancement algorithms, even the most sophisticated ones, do not improve speech intelligibility. Clearly, one reason is the fact that we often do not have a good estimate of the background noise spectrum, which is needed for the implementation of most algorithms. For that, accurate voice-activity detection algorithms are required. Much progress has been made in the design of voice-activity detection algorithms and noise-estimation algorithms (see review in [1, Ch. 9], some of which (e.g., [6]) are capable of continuously tracking, at least, the mean of the noise spectrum. Noise-estimation algorithms are known to perform well in stationary background noise (e.g., car) environments. Evidence of this was provided by Hu and Loizou [4] wherein a small improvement ($< 10\%$) in intelligibility was observed with speech processed in car environments, but not in other environments (e.g., babble). We believe that the small improvement was attributed to the stationarity of the car noise, which allowed for accurate noise estimation. This suggests that accurate noise estimation can contribute to improvement in intelligibility, but that alone cannot provide substantial improvements in intelligibility, since in practice we will never be able to track accurately the spectrum of nonstationary noise. For that reason, we believe that the absence of intelligibility improvement with existing speech enhancement algorithms is not entirely due to the lack of accurate estimates of the noise spectrum.

In this paper, we discuss other factors that are responsible for the absence of intelligibility improvement with existing algorithms. The majority of these factors center around the fact that none of the existing algorithms are designed to improve speech intelligibility, as they utilize a cost function that does not necessarily correlate with speech intelligibility. The statistical-model based algorithms (e.g., MMSE, Wiener filter), for instance, derive the magnitude spectra by minimizing the mean-squared error (MSE) between the clean and estimated (magnitude or power) spectra (e.g., [7]). The MSE metric, however, pays no attention to positive or negative differences between the clean and estimated spectra. A positive difference between the clean and estimated spectra would signify attenuation distortion, while a negative spectral difference would signify amplification distortion. The perceptual effect of these two distortions on speech intelligibility cannot be assumed to be equivalent. The subspace techniques (e.g., [8]) were designed to minimize a mathematically-derived speech distortion measure, but make no attempt

Manuscript received July 25, 2009; revised October 16, 2009; accepted January 26, 2010. Date of publication March 11, 2010; date of current version October 01, 2010. This work was supported by the NIDCD/NIH under Grant R01 DC007527. The associate editor coordinating the review of this manuscript and approving it for publication was Dr. Malcolm Slaney.

P. C. Loizou is with the Department of Electrical Engineering, University of Texas at Dallas, Richardson, TX 75083-0688 USA (e-mail: loizou@utdallas.edu).

G. Kim is with the School of Electronic Engineering, College of Information and Communication, Daegu University, Daegu 712-714, Korea (e-mail: imkgb27@gmail.com).

Color versions of one or more of the figures in this paper are available online at <http://ieeexplore.ieee.org>.

Digital Object Identifier 10.1109/TASL.2010.2045180

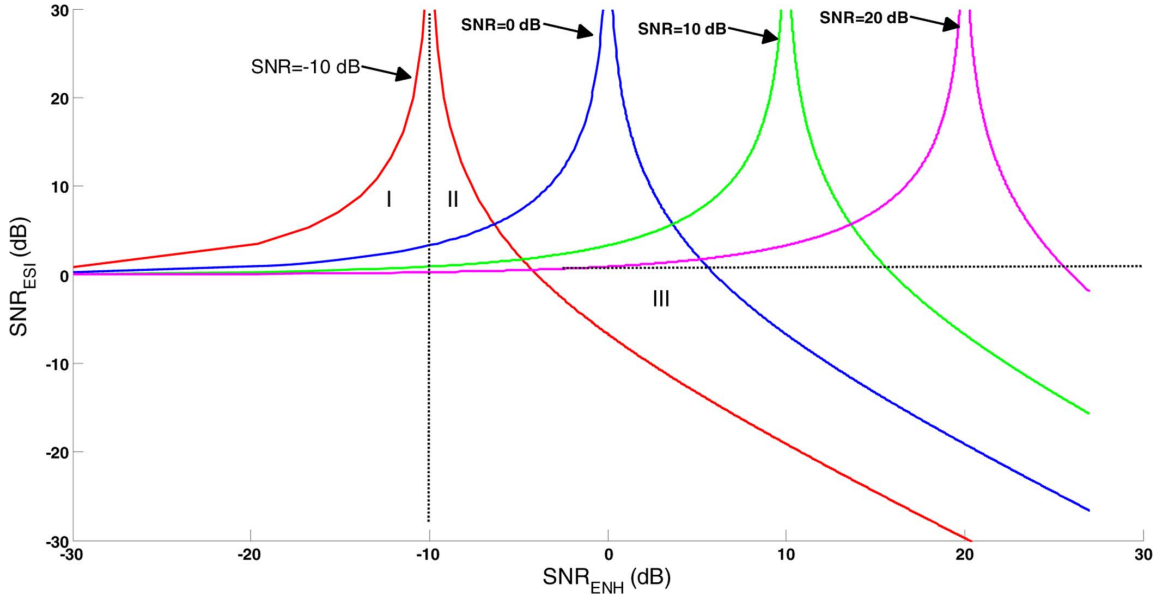


Fig. 1. Plot showing the relationship between SNR_{ESI} and SNR_{ENH} for fixed values of SNR.

to differentiate between the two aforementioned distortions. In this paper, we will show analytically that if we can somehow manage or control these two types of distortions, then we should expect to receive large gains in intelligibility. To further support our hypothesis, intelligibility listening tests are conducted with normal-hearing listeners.

II. IMPOSING CONSTRAINTS ON THE ESTIMATED MAGNITUDE SPECTRA

To gain a better understanding on the impact of the two distortions on speech intelligibility, we use an objective function that has been found to correlate highly ($r = 0.81$) with speech intelligibility [9]. This measure is the frequency-domain version of the well-known segmental SNR measure. The time-domain segmental (and overall) SNR measure has been used widely and frequently for evaluating speech quality in speech coding and enhancement applications [10], [11]. Results reported in the previous studies [9], [12], however, demonstrated that the time-domain SNR measure does not correlate highly with either quality or speech intelligibility. In contrast, the frequency domain version of the segmental SNR measure [13] has been shown to correlate highly with both speech quality and speech intelligibility. In the present study, we refer to this measure as the signal-to-residual spectrum measure SNR_{ESI} (defined below). The correlation of the SNR_{ESI} measure with speech intelligibility was found to be 0.81 [9] and the correlation with speech quality was found to be 0.85 [12]. The two main advantages in computing the SNR_{ESI} measure in the frequency domain include: 1) the use of critical-band frequency spacing for proper modeling of the frequency selectivity of normal-hearing listeners; 2) the use of perceptually motivated weighting functions which can be applied to individual bands [9]. The use of signal-dependent weighting functions in the computation of the SNR_{ESI} measure was found to be particularly necessary for predicting the intelligibility of speech corrupted by (fluctuating) nonstationary noise

[9]. We thus believe that it is the combination of these two attractive features in the computation of the SNR_{ESI} measure that contributes to its high correlation with speech intelligibility.

Let $\text{SNR}_{\text{ESI}}(k)$ denote the signal-to-residual spectrum ratio at frequency bin k

$$\text{SNR}_{\text{ESI}}(k) = \frac{X^2(k)}{(X(k) - \hat{X}(k))^2} \quad (1)$$

where $X(k)$ denotes the clean magnitude spectrum and $\hat{X}(k)$ denotes the magnitude spectrum *estimated* by a speech-enhancement algorithm. Dividing both numerator and denominator by $D^2(k)$, where $D(k)$ denotes the noise magnitude spectrum, we get

$$\text{SNR}_{\text{ESI}}(k) = \frac{\text{SNR}(k)}{(\sqrt{\text{SNR}(k)} - \sqrt{\text{SNR}_{\text{ENH}}(k)})^2} \quad (2)$$

where $\text{SNR}(k) \triangleq X^2(k)/D^2(k)$ is the true instantaneous SNR at bin k , and $\text{SNR}_{\text{ENH}}(k) \triangleq \hat{X}^2(k)/D^2(k)$ is the enhanced SNR.¹ Fig. 1 plots $\text{SNR}_{\text{ESI}}(k)$ as a function of $\text{SNR}_{\text{ENH}}(k)$, for fixed values of SNR. The singularity in the function stems from the fact that when $\text{SNR}(k) = \text{SNR}_{\text{ENH}}(k)$, $\text{SNR}_{\text{ESI}}(k) = \infty$. Fig. 1 provides important insights about the contributions of the two distortions to $\text{SNR}_{\text{ESI}}(k)$, and for convenience, we divide the figure into multiple regions according to the distortions introduced.

Region I. In this region, $\hat{X}(k) \leq X(k)$, suggesting only attenuation distortion.

Region II. In this region, $X(k) < \hat{X}(k) \leq 2 \cdot X(k)$ suggesting amplification distortion up to 6.02 dB.

Region III. In this region, $\hat{X}(k) > 2 \cdot X(k)$ suggesting amplification distortion of 6.02 dB or greater.

¹Note that the defined enhanced SNR_{ENH} is not the same as the output SNR, since the background noise is not processed separately by the enhancement algorithm.

TABLE I
PERCENTAGE OF FREQUENCY BINS FALLING IN THE THREE REGIONS AFTER PROCESSING NOISY SPEECH BY THE THREE ENHANCEMENT ALGORITHMS

SNR Level	Algorithm	Region I	Region II	Region III
0 dB	Wiener	45.33 %	14.64 %	40.03 %
-5 dB		37.71 %	12.71 %	49.58 %
0 dB	RDC_mag	46.86 %	15.45 %	37.69 %
-5 dB		35.88 %	14.49 %	49.63 %
0 dB	RDC_power	36.81 %	18.14 %	45.05 %
-5 dB		28.08 %	14.88 %	57.04 %

From the above, we can deduce that in the union of Regions I and II, which we denote as Region I + II, we have the following constraint:

$$\hat{X}(k) \leq 2 \cdot X(k). \quad (3)$$

The constraint in Region I stems from the fact that in this region, $\text{SNR}_{\text{ENH}} \leq \text{SNR}$ leading to $\hat{X}(k) \leq X(k)$. The constraint in Region II stems from the fact that in this region $\text{SNR} \leq \text{SNR}_{\text{ENH}} \leq \text{SNR} + 6.02$ dB. Finally, the condition in Region III stems from the fact that in this region $\text{SNR}_{\text{ENH}} \geq \text{SNR} + 6.02$ dB. It is clear from the above definitions of these three regions that in order to maximize in some respect the SNR_{ESI} (and consequently maximize speech intelligibility), the estimated magnitude spectra $\hat{X}(k)$ need to be contained in regions I and II (note that the trivial, but not useful, solution that maximizes SNR_{ESI} is $\hat{X} = X(k)$). Intelligibility listening tests were conducted to test this hypothesis. If the hypothesis holds, then we expect to see large improvements in intelligibility.

It is reasonable to ask how often the above distortions occur when corrupted speech is processed by conventional speech-enhancement algorithms. To answer this question, we tabulate in Table I the frequency of occurrences of the two distortions for speech processed by three different (but commonly used) algorithms at two different SNR levels. Table I provides the average percentage of frequency bins falling in each of the three regions. To compute, for instance, the percentage of bins falling in Region I we counted the number of bins satisfying the constraint in Region I, and divided that by the total number of frequency bins, as determined by the size of the discrete Fourier transform (DFT). This was done at each frame after processing corrupted speech with an enhancement algorithm, and averaging the percentages over all frames in a sentence. As can be seen, nearly half of the bins fall in Region I which is characterized by attenuation distortion, while the other half of the bins fall in Region III, which is characterized by amplification distortion in excess of 6.02 dB. A small percentage (12%–18%) of bins was found to fall in Region II which is characterized by low amplification distortion, less than 6.02 dB. The perceptual consequences of the two distortions on speech intelligibility are not clear. For one, it is not clear which of the two distortions has the most detrimental effect on speech intelligibility. Listening tests are conducted to provide answers to these questions, and these tests are described next.

III. INTELLIGIBILITY LISTENING TESTS

Algorithms Tested

The noise-corrupted sentences were processed by three different speech enhancement algorithms that included the Wiener

algorithm based on *a priori* SNR estimation [14] and two spectral-subtractive algorithms based on reduced delay convolution [15]. The sentences were segmented into overlapping segments of 160 samples (20 ms) with 50% overlap. Each segment was Hann windowed and transformed using a 160-point discrete Fourier transform (DFT). Let $Y(k, t)$ denote the magnitude of the noisy spectrum at time frame t and frequency bin k . Then, the estimate of the signal spectrum magnitude is obtained by multiplying $Y(k, t)$ with a gain function $G(k, t)$ as follows:

$$\hat{X}(k, t) = G(k, t) \cdot Y(k, t). \quad (4)$$

Three different gain functions were considered in the present study. The Wiener gain function is based on the *a priori* SNR and is given by

$$G_{\text{Wiener}}(k, t) = \sqrt{\frac{\text{SNR}_{\text{prio}}(k, t)}{1 + \text{SNR}_{\text{prio}}(k, t)}} \quad (5)$$

where SNR_{prio} is the *a priori* SNR estimated using the decision-directed approach as follows:

$$\text{SNR}_{\text{prio}}(k, t) = \alpha \cdot \frac{\hat{X}^2(k, t-1)}{\hat{P}_D^2(k, t-1)} + (1 - \alpha) \cdot \max \left[\frac{Y^2(k, t)}{\hat{P}_D^2(k, t)} - 1, 0 \right] \quad (6)$$

where $\hat{P}_D^2(k, t)$ is the estimate of the power spectral density of background noise and α is a smoothing constant (typically set to $\alpha = 0.98$). The spectral subtractive algorithms are based on reduced delay convolution [15] and the gain functions for magnitude subtraction and power subtraction are given respectively by

$$G_{\text{RDC_mag}}(k, t) = \max \left[0, 1 - \frac{\beta}{\sqrt{\text{SNR}_{\text{post}}(k, t)}} \right] \quad (7)$$

$$G_{\text{RDC_pow}}(k, t) = \sqrt{\max \left[0, 1 - \frac{\beta}{\text{SNR}_{\text{post}}(k, t)} \right]} \quad (8)$$

where $\text{SNR}_{\text{post}}(k, t) \triangleq \hat{P}_Y^2(k, t) / \hat{P}_D^2(k, t)$ denotes the *a posteriori* SNR, and $\hat{P}_Y^2(k, t)$ denotes the estimate of the noisy speech power spectral density computed using the periodogram method, and β is the subtraction factor which was set to $\beta = 0.7$ as per [15]. We denote the magnitude spectral-subtraction algorithm as RDC_mag and the power spectral-subtraction algorithm as RDC_pow.

The three gain functions examined are plotted in Fig. 2. As can be seen, the three algorithms differ in the shape of their gain

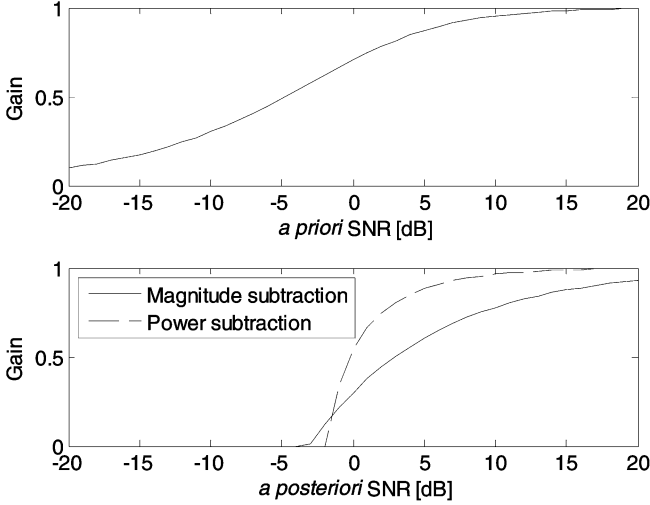


Fig. 2. Suppression curves of the Wiener filtering algorithm (top panel) and two spectral-subtractive algorithms (bottom panel).

functions. The Wiener gain function is the least aggressive, in terms of suppression, providing small attenuation even at extremely low SNR levels, while the RDC_mag algorithm is the most aggressive eliminating spectral components at extremely low SNR levels. The three gain functions span a wide range of suppression options, which is one of the reasons for selecting them. We will thus be able to test our hypotheses about the effect of the constraints on speech intelligibility with algorithms encompassing a wide range of suppression varying from aggressive to least aggressive. The RDC_mag algorithm, in particular, was chosen because it performed poorly in terms of speech intelligibility [4]. The intelligibility of speech processed by the RDC_mag algorithm was found in several noisy conditions to be significantly lower than that obtained with unprocessed (noise corrupted) speech. We will thus examine whether it is possible to obtain improvement in intelligibility with the proposed constraints, even in scenarios where the enhancement algorithm (e.g., RDC_mag) is known to perform poorly relative to unprocessed speech.

Oracle experiments were run in order to assess the full potential on speech intelligibility when the proposed constraints are implemented. We thus assumed knowledge of the magnitude spectrum of the clean speech signal. The various constraints were implemented as follows. The noisy speech signal was first segmented into 20 ms frames (with 50% overlap between frames), and then processed through one of the three enhancement algorithms, producing at each frame the estimated magnitude spectrum $\hat{X}(k)$. The noise estimation algorithm proposed by Rangachari and Loizou [16] was used for estimating the noise spectrum in (6)–(8). The estimated magnitude spectrum $\hat{X}(k)$ was compared against the true spectrum $X(k)$, and spectrum components satisfying the constraint were retained, while spectral components violating the constraints were zeroed-out. For the implementation of the Region I constraint, for instance, the modified magnitude spectrum $X_M(k)$ was computed as follows:

$$X_M(k) = \begin{cases} \hat{X}(k), & \text{if } \hat{X}(k) < X(k) \\ 0, & \text{else.} \end{cases} \quad (9)$$

An inverse discrete Fourier transform (IDFT) was finally taken of $X_M(k)$ (using the noisy speech signal's phase spectrum) to reconstruct the time-domain signal. The overlap-and-add technique was subsequently used to synthesize the signal. As shown in (9), the constraints are implemented by applying a binary mask to the *estimated* magnitude spectrum (more on this later).

Fig. 3 shows example spectrograms of signals synthesized using the Region I constraints [panel (d)]. The original signal was corrupted with babble at -5 dB SNR. The Wiener algorithm was used in this example, and speech processed with the Wiener algorithm is shown in panel (c). As can be seen in panel (d), the signal processed using the Region I constraints resembles the clean signal, with most of the residual noise removed and the consonant onsets /offsets made clear.

A. Methods and Procedure

Seven normal-hearing listeners participated in the listening experiments, and all listeners were paid for their participation. The listeners participated in a total of 32 conditions ($=2$ SNR levels (-5 dB, 0 dB) $\times 16 (= 3 \times 5 + 1)$ processing conditions). For each SNR level, the processing conditions included speech processed using three different speech enhancement (SE) algorithms with 1) no constraints imposed, 2) Region I constraints, 3) Region II constraints, 4) Region I + II constraints, and 5) Region III constraints. For comparative purposes, subjects were also presented with noise-corrupted (unprocessed) stimuli.

The listening experiment was performed in a sound-proof room (Acoustic Systems, Inc.) using a PC connected to a Tucker-Davis system 3. Stimuli were played to the listeners monaurally through Sennheiser HD 250 Linear II circumaural headphones at a comfortable listening level. Prior to the sentence test, each subject listened to a set of noise-corrupted sentences to be familiarized with the testing procedure. During the test, subjects were asked to write down the words they heard. Two lists of sentences (i.e., 20 sentences) were selected from the IEEE database [17] and used for each condition, with none of the lists repeated across conditions. The order of the conditions was randomized across subjects. The testing session lasted for about 2 h. Five-minute breaks were given to the subjects every 30 minutes.

Sentences taken from the IEEE database [17] were used for test material.² The sentences in the IEEE database are phonetically balanced with relatively low word-context predictability. The sentences were originally recorded at a sampling rate of 25 kHz and downsampled to 8 kHz (the recordings are available from a CD accompanying the book in [1]). Noisy speech was generated by adding babble noise at 0 dB and -5 dB SNR. The babble noise was produced by 20 talkers with equal number

²A sentence recognition test was chosen over a diagnostic rhyme test [18] for assessment of intelligibility for several reasons. Sentence tests: 1) better reflect real-world communicative situations, 2) are open-set tests, and as such scores may vary from a low of 0% correct to 100% correct (in contrast, the DRT test is a closed-set test, has a chance score of 50% and needs to be corrected for chance). The sentence materials (IEEE corpus) chosen contain contextual information, however, that information is controlled by design. The IEEE corpus contains phonetically balanced sentences and is organized into lists of ten sentences each. All sentence lists were designed to be equally intelligible, thereby allowing us to assess speech intelligibility in different conditions without being concerned that a particular list is more intelligible than another.

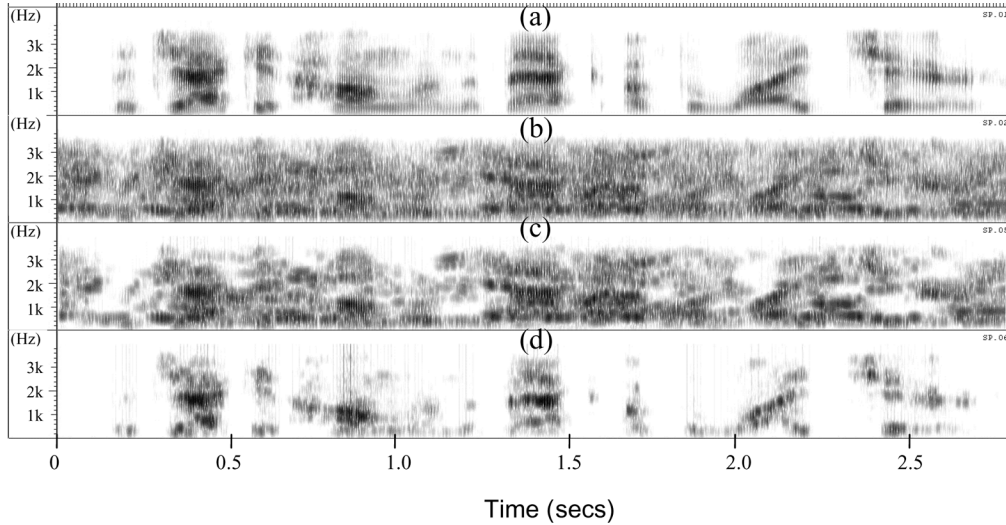


Fig. 3. Wide-band spectrograms of the clean signal [panel (a)], noisy signal in -5 dB SNR babble [panel (b)], signal processed by the Wiener algorithm [panel (c)], and signal processed by the Wiener algorithm after imposing the constraints in Region I [panel (d)].

of female and male talkers. To simulate the receiving frequency characteristics of telephone handsets, the speech and noise signals were filtered by the modified intermediate reference system (IRS) filters used in ITU-T P.862 [19]. Telephone speech was used as it is considered particularly challenging (in terms of intelligibility) owing to its limited bandwidth (300–3200 Hz). Consequently, we did not expect the performance to be limited by ceiling effects.

B. Results

Fig. 4 shows the results of the listening tests expressed in terms of the percentage of words identified correctly by normal-hearing listeners. The bars indicated as “UN” show the scores obtained with noise-corrupted (unprocessed) stimuli, while the bars indicated as “SE” show the baseline scores obtained with the three enhancement algorithms (no constraints imposed). As shown in Fig. 4, performance improved dramatically when the Region I constraints were imposed. Consistent improvement in intelligibility was obtained with the Region I constraints for all three speech-enhancement algorithms examined. Performance at -5 dB SNR with the Wiener algorithm, for instance, improved from 10% correct when no constraints were imposed, to 90% correct when Region I constraints were imposed. Substantial improvements in intelligibility were also noted for the two spectral-subtractive algorithms examined. Performance with Region II constraints seemed to be dependent on the speech-enhancement algorithm used, with good performance obtained with the Wiener algorithm, and poor performance obtained with the two spectral-subtractive algorithms. Large improvements in intelligibility were obtained with Region I + II constraints for all three algorithms tested and for both SNR levels. Finally, performance degraded to near zero when Region III constraints were imposed for all three algorithms tested and for both SNR levels.

Statistical tests, based on Fisher’s LSD test, were run to assess significant differences between the scores obtained in the

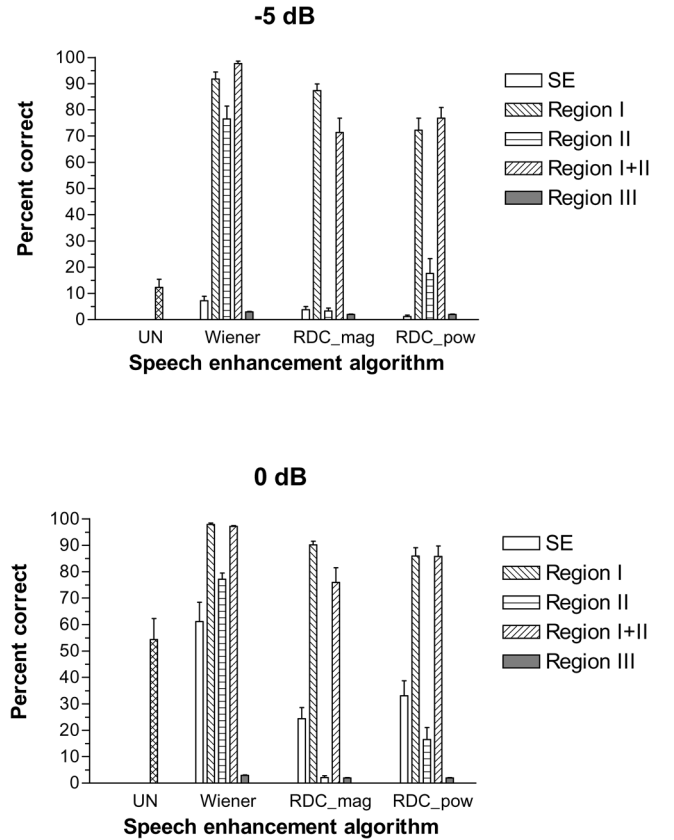


Fig. 4. Results, expressed in percentage of words identified correctly, from the intelligibility studies with human listeners. The bars indicated as “UN” show the scores obtained with noise-corrupted (unprocessed) stimuli, while the bars indicated as “SE” show the baseline scores obtained with the three enhancement algorithms (no constraints imposed). The intelligibility scores obtained with speech processed by the three enhancement algorithms after imposing four different constraints are labeled accordingly.

various constraint conditions. Performance of the Wiener algorithm with Region I constraints did not differ statistically ($p > 0.05$) from performance obtained with the Region I + II constraints. Similarly, performance of the RDC_pow algorithm with Region I constraints did not differ statistically ($p > 0.05$)

from performance obtained with the Region I + II constraints. This was found to be true for both SNR levels and for both Wiener and RDC_pow algorithms. In contrast, performance obtained with the RDC_mag algorithm with Region I constraints was significantly higher ($p < 0.05$) than performance obtained with Region I + II constraints. Performance obtained with the Wiener algorithm (with no constraints) did not differ significantly ($p > 0.05$) from performance obtained with unprocessed (noise corrupted) sentences for both SNR levels tested. Performance obtained at -5 dB SNR with the two spectral-subtractive algorithms did not differ significantly ($p > 0.05$) from performance obtained with unprocessed (noise corrupted) sentences, but was found to be significantly ($p < 0.05$) lower than performance with unprocessed sentences at 0 dB SNR. The latter outcome is consistent with the findings reported by Hu and Loizou [4].

In summary, the above analysis indicates that the Region I and Region I + II constraints are the most robust in terms of yielding consistently large benefits in intelligibility *independent* of the speech-enhancement algorithm used. Substantial improvements in intelligibility (85 percentage points at -5 dB SNR and nearly 70 percentage points at 0 dB SNR) were obtained even with the RDC_mag algorithm, which was found in our previous study [4], as well as in the present study, to degrade speech intelligibility in some noisy conditions. Of the three enhancement algorithms examined, the Wiener algorithm is recommended when imposing Region I or Region I + II constraints, as this algorithm yielded the largest gains in intelligibility for both SNR levels tested. Based on data from Table I, there does not seem to be a correlation between the numbers of frequency bins falling in the three regions with speech intelligibility gains. The RDC_pow algorithm, for instance, yielded roughly the same number of frequency bins in Region I as the Wiener filtering algorithm, yet the latter algorithm obtained larger improvements in intelligibility. We attribute the difference in performance to the shape of the suppression function.

A difference in outcomes in Region II was observed between the Wiener and spectral subtractive algorithms. Compared to the performance obtained by the subtractive algorithms in Region II, the performance of the Wiener algorithm was substantially higher. To analyze this, we examined the frequency dependence of the distortions in Region II. More precisely, we examined whether distortions in region II (as introduced by the three different algorithms) occurred more frequently within a specific frequency region. We first divided the signal bandwidth into three frequency regions: low-frequency ($0-1$ kHz), mid-frequency ($1-2$ kHz) and high-frequency regions ($2-4$ kHz). We then computed the percentage of bins falling in each of the three frequency regions for speech processed by the three algorithms (only accounting for distortions in Region II). The results, averaged over 20 sentences, are shown in Fig. 5. As can be seen from this Figure, a slightly higher percentage of bins were observed in the lower frequency region ($0-1$ kHz) for the Wiener algorithm compared to the spectral subtractive algorithms. The higher percentage in the lower frequency region ($0-1$ kHz), where the first

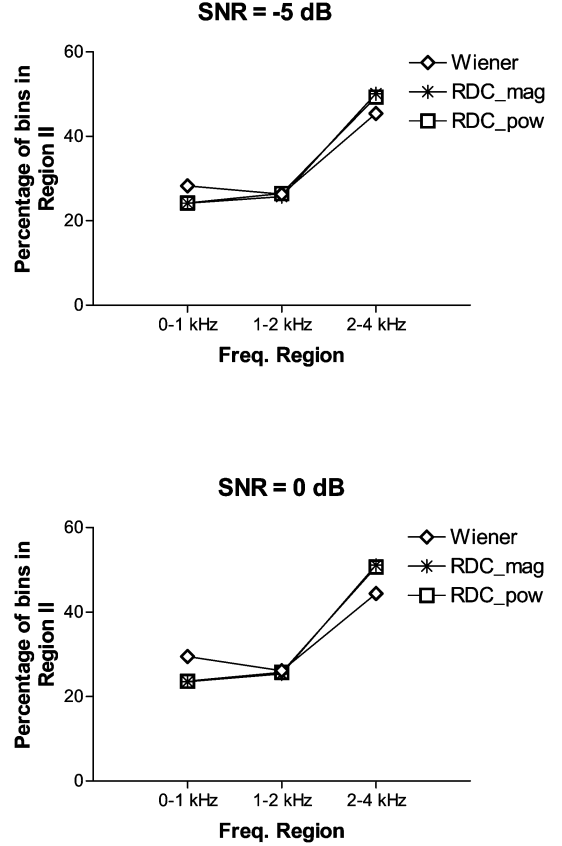


Fig. 5. Percentage of bins falling in three different frequency regions (Region II constraints).

formant frequency resides, might partially explain the better intelligibility scores, but this difference was rather small and not enough to account for the difference in intelligibility in Region II between the Wiener algorithm and the spectral subtractive algorithms.

We continued the analysis of Region II, and considered computing the histograms of the following estimation error: $\hat{X}_{dB} - X_{dB}$, where the subscript indicates that the magnitudes are expressed in dB. Note that this error is always positive and is upper bounded by 6.02 dB in Region II. The resulting histograms are shown in Fig. 6. As can be seen from this figure, magnitude errors smaller than 1 dB were made more frequently by the Wiener filtering algorithm for both SNR conditions, compared to the uniformly-distributed errors (at least at -5 dB SNR) made by the spectral-subtractive algorithms. This suggests that the Wiener filtering algorithm correctly estimates the true magnitude spectra more often compared to the subtractive algorithms, at least in Region II. We believe that this could be the reason that the Wiener algorithm performed better than the subtractive algorithms in Region II.

Performance in Region III (Fig. 4) was extremely low (near 0% correct) for all three algorithms tested. We believe that this was due to the excess masking of the target signal in this region. Amplification distortions in excess of 6.02 dB were introduced. In Region III, the masker overpowered the target signal, rendering it unintelligible.

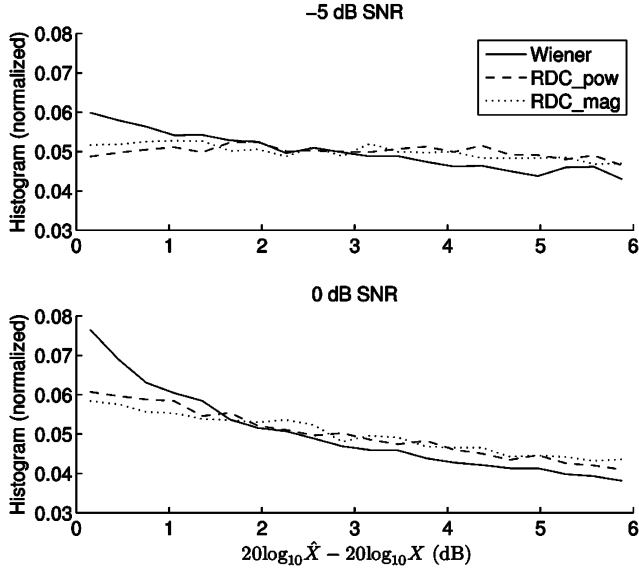


Fig. 6. Normalized histograms (probability mass functions) of the difference between the estimated and clean speech magnitudes in Region II.

IV. RELATIONSHIP BETWEEN PROPOSED RESIDUAL CONSTRAINTS AND THE IDEAL BINARY MASK

As shown in (9), the modified spectrum (with the proposed constraints incorporated) can be obtained by applying a binary mask to the enhanced spectrum. In computational auditory scene analysis (CASA) applications, a binary mask is often applied to the noisy speech spectrum to recover the target signal [20]–[23]. In this section, we show that there exists a relationship between the proposed residual constraints (and associated binary mask) and the ideal binary mask used in CASA and robust speech recognition applications (e.g., [21]). The goal of CASA techniques is to segregate the target signal from the sound mixtures, and several techniques have been proposed in the literature to achieve that [23]. These techniques can be model-based [24], [25] or based on auditory scene analysis principles [26]. Some of the latter techniques use the ideal time–frequency (T-F) binary mask [20], [21], [27]. The ideal binary “mask” (IdBM) takes values of zero or one, and is constructed by comparing the local SNR in each T-F unit (or frequency bin) against a threshold (e.g., 0 dB). It is commonly applied to the T-F representation of a mixture signal and eliminates portions of a signal (those assigned to a “zero” value) while allowing others (those assigned to a “one” value) to pass through intact. The ideal binary mask provides the only known criterion ($\text{SNR} \geq \delta$ dB, for a preset threshold δ) for improving speech intelligibility, and this was confirmed by several intelligibility studies with normal-hearing [28], [29] and hearing-impaired listeners [30], [31]. IdBM techniques often introduce musical noise, caused by errors in the estimation of the time–frequency masks and manifested in isolated T-F units. A number of techniques have been proposed to suppress musical noise distortions introduced by IdBM techniques [32], [33]. While musical noise might be distracting to the listeners, it has not been found to be detrimental in terms of speech intelligibility. This was confirmed in two listening studies with IdBM-processed speech [28], [29] and in one study with

estimated time–frequency masks [34]. Despite the presence of musical noise, normal-hearing listeners were able to recognize estimated [34] and ideal binary-masked [28], [29] speech with nearly 100% accuracy.

The reasons for the improvement in intelligibility with IdBM are not very clear. Li and Wang [35] argued that the IdBM maximizes the SNR as it minimizes the sum of missing target energy that is discarded and the masker energy that is retained. More specifically, it was proven that the IdBM criterion maximizes the SNR_{ESI} metric given in (1) [35]. The IdBM was also shown to maximize the time-domain based segmental and overall SNR measures, which are often used for assessment of speech quality. Neither of these measures, however, correlates with speech intelligibility [9]. We provide proof in the Appendix that the IdBM criterion maximizes the geometric average of the spectral SNRs, and subsequently maximizes the articulation index (AI), a metric known to correlate highly with speech intelligibility [36].

As it turns out, the ideal binary mask is not only related to the proposed residual constraints, but is also a special case of the proposed residual constraint for regions I and II. Put differently, the proposed binary mask [see example in (9)] is a generalized form of the ideal binary mask used in CASA applications. As mentioned earlier, if the estimated magnitude spectrum is restricted to fall within regions I and II, then the SNR_{ESI} metric will always be greater than 0 dB. Hence, imposing constraints in region I + II ensures that SNR_{ESI} is always positive and greater than 1 (i.e., > 0 dB). As demonstrated in Fig. 4, the stimuli constrained in region I + II consistently improved speech intelligibility for all three enhancement algorithms tested. As mentioned earlier, the composite constraint required for the estimated magnitude spectra to fall in region I + II is given by

$$\hat{X}(k) \leq 2 \cdot X(k) \quad (10)$$

which after squaring both sides becomes

$$\hat{X}^2(k) \leq 4 \cdot X^2(k). \quad (11)$$

If we now assume that $\hat{X}(k) = Y(k)$, i.e., that the noisy signal is not processed by an enhancement algorithm, then $\hat{X}^2(k) = Y^2(k) = X^2(k) + D^2(k)$, and (11) reduces to

$$\begin{aligned} X^2(k) &\geq \frac{1}{3} D^2(k) \\ \text{SNR}(k) &\geq \frac{1}{3}. \end{aligned} \quad (12)$$

In dB, the above equation suggests that the SNR needs to be larger than a threshold of -4.77 dB. Equation (12) is nothing but the criterion used in the construction of the ideal binary mask. The only difference is that the threshold used is -4.77 dB, rather than 0 or -3 dB, which are most often used in applications of the IdBM [27]. In terms of obtaining intelligibility improvement, however, either threshold is acceptable. The previous intelligibility studies confirmed that there exists a plateau in performance when intelligibility was measured as a function of the SNR threshold [28], [29], [37]. In the study conducted by Li and Loizou [37], for instance, the plateau in performance

(nearly 100% correct) ranged from an SNR threshold of -20 dB to 0 dB.

As shown in Fig. 4, the constraint stated in (10) guarantees substantial improvement in intelligibility for all three algorithms tested. The ideal binary mask is a special case of this constraint when no enhancement algorithm is used, i.e., when no processing is applied to the noisy speech signal. Unlike the criterion used in the binary mask [(12)], the proposed constraints [(10)] do not involve the noise spectrum, at least explicitly. In contrast, the ideal binary mask criterion requires access to the true noise spectrum, which is extremely challenging to obtain at very low SNR levels (e.g., $\text{SNR} < 0$ dB). Attempts to estimate the binary mask using existing speech enhancement algorithms met with limited success (e.g., [38] and [39]), and performance, in terms of detection rates was found to be relatively poor. It remains to be seen whether it is easier to estimate the proposed binary mask [e.g., (9)], given that it does not require access to the true noise spectrum.

V. DISCUSSION AND CONCLUSION

Current speech enhancement algorithms can improve speech quality but not speech intelligibility [4]. Quality and intelligibility are two of the many attributes (or dimensions) of speech and the two are not necessarily equivalent. Hu and Loizou [2], [4] showed that algorithms that improve speech quality do not improve speech intelligibility. The subspace algorithm, for instance, was found to perform the worst in terms of overall quality [2], but performed well in terms of preserving speech intelligibility [4]. In fact, in babble noise (0 dB SNR), the subspace algorithm performed significantly better than the logMMSE algorithm [40], which was found to be among the algorithms yielding the highest overall speech quality [2].

The findings of the present study suggest two interrelated reasons for the absence of intelligibility improvement with existing speech enhancement (SE) algorithms. First, and foremost, SE algorithms do not pay attention to the two types of distortions introduced when applying the suppression function to noisy speech spectra. Both distortions are treated equally in most SE algorithms, since the MSE metric is used in the derivation of most suppression functions (e.g., [7]). As demonstrated in Fig. 4, however, the perceptual effects of the two distortions on speech intelligibility are not equal. Of the two types of distortion, the amplification distortion (in excess of 6.02 dB) was found to bear the most detrimental effect on speech intelligibility (see Fig. 4). Performance dropped near zero when stimuli were constrained in region III. Theoretically, we believe that this is so because this type of distortion (region III) leads to negative values of SNR_{ESI} (see Fig. 1). In contrast, the attenuation distortion (region I) was found to yield the least effect on intelligibility. In fact, when the region I constraint was imposed, large gains in intelligibility were realized. Performance at -5 dB SNR, improved from 5% correct with stimuli enhanced with the Wiener algorithm to 90% correct when region I constraint was imposed. Theoretically, we believe that the improvement in intelligibility is due to the fact that region I always ensures that $\text{SNR}_{\text{ESI}} \geq 0$ dB. Maximizing SNR_{ESI} ought to maximize intelligibility, given the high correlation of a weighted-version of SNR_{ESI} (termed $\text{fwSNR}_{\text{seg}}$ [9], [11]) with speech intelligibility. Hence, by imposing the appropriate constraints [see

(10)], we can ensure that $\text{SNR}_{\text{ESI}} \geq 0$ dB, and subsequently obtain large gains in intelligibility.

Second, none of the existing SE algorithms was designed to maximize a metric that correlates highly with intelligibility. The only known metric, which is widely used in CASA, is the ideal binary mask criterion. We provided a proof in the Appendix that this metric maximizes the articulation index, an index that is known to correlate highly with speech intelligibility [36]. Hence, it is not surprising that speech synthesized based on the IdBM criterion improves intelligibility [28], [29], [37]. In fact, it restores speech intelligibility to the level attained in quiet (near 100% correct) even for sentences corrupted by background noise at SNR levels as low as -10 dB SNR [29]. As shown in previous section, the IdBM criterion is a special case of the proposed constraint in region I + II, when no suppression function is applied to the noisy spectra, i.e., when $\hat{X}(k) = Y(k)$.

In summary, in order for SE algorithms to improve speech intelligibility they need to treat the two types of distortions differently. More specifically, SE algorithms need to be designed so as to minimize the amplification distortions. As the data in Fig. 4 demonstrated, even spectral-subtractive algorithms can improve speech intelligibility if the amplification distortions are properly controlled. In practice, the proposed constraints can be imposed and incorporated in the derivation of the noise suppression function. That is, rather than focusing on minimizing a squared-error criterion (as done in the derivation of MMSE algorithms), we can focus instead on minimizing a given criterion *subject to* the proposed constraints. The speech enhancement problem is thus converted to a constrained minimization problem. Alternatively, and perhaps, equivalently, SE algorithms need to be designed so as to maximize a metric (e.g., SNR_{ESI} , AI) that is known to correlate highly with speech intelligibility (for a review of such metrics, see [9]). For instance, SE algorithms need to be designed to maximize SNR_{ESI} rather than minimize an unconstrained (mean) squared-error cost function, as done by most statistical-model based algorithms (e.g., [7]). Algorithms that maximize the SNR_{ESI} metric are likely to provide substantial gains in intelligibility.

APPENDIX

In this Appendix, we provide analytical proof that the IdBM criterion is optimal in that it maximizes the geometric average of the spectral SNRs. We also show that maximizing the geometric average of SNRs is equivalent to maximizing a simplified form of the articulation index (AI),³ an objective measure used for predicting speech intelligibility [36], [44].

³The AI index has been shown to predict reliably speech intelligibility by normal-hearing [36] and hearing-impaired listeners [41] (the refined AI index is known as the speech intelligibility index and is documented in [42]). The AI measure, however, has a few limitations. First, the AI measure has been validated for the most part only for stationary masking noise since it is based on the long-term average spectra, computed over 125 ms intervals, of the speech and masker signals [42]. As such, it cannot be applied to situations in which speech is embedded in fluctuating maskers e.g., competing talkers. Several attempts have been made, however, to extend the AI measure to assess speech intelligibility in fluctuating maskers (e.g., see [9], [43]). Second, the AI measure cannot predict synergistic effects as evident in the perception of disjoint frequency bands. This is so due to the assumption that individual frequency bands contribute independently to AI.

Consider the following weighted (geometric) average of SNRs computed across N frequency bins

$$F = \sum_{j=1}^N I_j \cdot \text{SNR}(j). \quad (\text{A.1})$$

where $\text{SNR}(j) = 10 \log_{10} (X^2(j)/D^2(j))$ is the SNR in bin (or channel) j and I_j are the weights ($0 \leq I_j \leq 1$) applied to each frequency bin. We consider the following question: how should the weights I_j be chosen such that the overall SNR (i.e., F) given in (A.1) is maximized? The rationale for wanting to maximize F stems from the fact that F is similar to the articulation index (more on this below). The optimal weights that maximize F in (1) are given by

$$I_j^* = \begin{cases} 1, & \text{if } \text{SNR}(j) > 0 \\ 0 & \text{if } \text{SNR}(j) \leq 0 \end{cases} \quad (\text{A.2})$$

which is no other than the IdBM criterion. To see why the weights given in (A.2) are optimal, we can consider two extreme cases in which either $\text{SNR}(j) \leq 0$ or $\text{SNR}(j) \geq 0$ in all frequency bins. If $\text{SNR}(j) \leq 0$ in all bins, then we have the following upper bound on the value of F

$$\sum_{j=1}^N I_j \cdot \text{SNR}(j) \leq 0. \quad (\text{A.3})$$

Similarly, if $\text{SNR}(j) \geq 0$ in all bins, then we have the following upper bound:

$$\sum_{j=1}^N I_j \cdot \text{SNR}(j) \leq \sum_{j=1}^N \text{SNR}(j) \quad (\text{A.4})$$

Both upper bounds [maxima in (A.3) and (A.4)] are attained with the optimal weights given in (A.2). That is, the maximum in (A.3) is attained with $I_j = 0$ (for all j), while the maximum in (A.4) is attained with $I_j = 1$ (for all j).

It is important to note that the function F given in (A.1), is very similar to the articulation index defined by [42], [45], and [46]

$$AI = \sum_{j=1}^M W_j \cdot (\alpha \cdot \overline{\text{SNR}}(j) + \beta) \quad (\text{A.5})$$

where W_j are the band-importance functions ($0 \leq W_j \leq 1$), $\overline{\text{SNR}}(j)$ are the SNR values limited to the range of $[-15, 15]$ dB, M is the number of critical-bands, and α, β are constants ($\alpha = 1/30, \beta = 0.5$) used to ensure that the SNR is mapped within the range of $[0, 1]$. Maximization of AI in (A.5) will yield a similar optimal solution for the weights W_j as shown in (A.2), with the only difference being the SNR threshold (i.e., it will no longer be 0 dB). The AI assumes a value of 0 when the speech is completely masked and a value between 0 and 1 for SNRs ranging from -15 to 15 dB. In the original AI calculation [44], the band-importance functions W_j are fixed and their values depend on the type of speech material used. In our case, the importance functions W_j are not fixed, but are chosen dynamically according to (A.2) so as to maximize the geometric average of all SNRs across the spectrum. Hence, the main motivation behind maximizing F in (A.1) is to maximize the articulation index

(A.5), and consequently maximize the amount of retained information contributing to speech intelligibility. Hence, the weights W_j in (A.2) used in the construction of the ideal binary mask can be viewed as the optimal band-importance function W_j needed to maximize the simplified form of articulation index in (A.1). It is for this reason that we believe that the use of the IdBM criterion (A.2) always improves speech intelligibility [29].

REFERENCES

- [1] P. Loizou, *Speech Enhancement: Theory and Practice*. Boca Raton, FL: CRC, 2007.
- [2] Y. Hu and P. Loizou, "Subjective comparison and evaluation of speech enhancement algorithms," *Speech Commun.*, vol. 49, pp. 588–601, 2007.
- [3] J. Lim, "Evaluation of a correlation subtraction method for enhancing speech degraded by additive noise," *IEEE Trans. Acoust., Speech, Signal Process.*, vol. ASSP-37, no. 6, pp. 471–472, Dec. 1978.
- [4] Y. Hu and P. Loizou, "A comparative intelligibility study of single-microphone noise reduction algorithms," *J. Acoust. Soc. Amer.*, vol. 122, no. 3, pp. 1777–786, 2007.
- [5] R. Bentler, H. Wu, J. Kettel, and R. Hurtig, "Digital noise reduction: Outcomes from laboratory and field studies," *Int. J. Audiol.*, vol. 47, no. 8, pp. 447–460, 2008.
- [6] R. Martin, "Noise power spectral density estimation based on optimal smoothing and minimum statistics," *IEEE Trans. Speech Audio Process.*, vol. 9, no. 5, pp. 504–512, Jul. 2001.
- [7] Y. Ephraim and D. Malah, "Speech enhancement using a minimum mean-square error short-time spectral amplitude estimator," *IEEE Trans. Acoust., Speech, Signal Process.*, vol. ASSP-32, no. 6, pp. 1109–1121, Dec. 1984.
- [8] Y. Hu and P. Loizou, "A generalized subspace approach for enhancing speech corrupted by colored noise," *IEEE Trans. Speech Audio Process.*, vol. 11, no. 4, pp. 334–341, Jul. 2003.
- [9] J. Ma, Y. Hu, and P. Loizou, "Objective measures for predicting speech intelligibility in noisy conditions based on new band-importance functions," *J. Acoust. Soc. Amer.*, vol. 125, no. 5, pp. 3387–3405, 2009.
- [10] V. Grancharov and W. Kleijn, "Speech quality assessment," in *Handbook of Speech Processing*, J. Benesty, M. Sondhi, and Y. Huang, Eds. Berlin, Germany: Springer Verlag, 2008, pp. 83–99.
- [11] S. Quackenbush, T. Barnwell, and M. Clements, *Objective Measures of Speech Quality*. Englewood Cliffs, NJ: Prentice-Hall, 1988.
- [12] Y. Hu and P. Loizou, "Evaluation of objective quality measures for speech enhancement," *IEEE Trans. Audio Speech Lang Process.*, vol. 16, no. 1, pp. 229–238, Jan. 2008.
- [13] J. Tribolet, P. Noll, B. McDermott, and R. E. Crochiere, "A study of complexity and quality of speech waveform coders," in *Proc. IEEE Int. Conf. Acoust., Speech, Signal Process.*, 1978, pp. 586–590.
- [14] P. Scalart and J. Filho, "Speech enhancement based on a priori signal to noise estimation," in *Proc. IEEE Int. Conf. Acoust., Speech, Signal Process.*, 1996, pp. 629–632.
- [15] H. Gustafsson, S. Nordholm, and I. Claesson, "Spectral subtraction using reduced delay convolution and adaptive averaging," *IEEE Trans. Speech Audio Process.*, vol. 9, no. 8, pp. 799–807, Nov. 2001.
- [16] S. Rangachari and P. Loizou, "A noise-estimation algorithm for highly non-stationary environments," in *Speech Commun.*, 2006, vol. 48, pp. 220–231.
- [17] "IEEE recommended practice for speech quality measurements," *IEEE Trans. Audio Electroacoust.*, vol. AU-17, no. 3, pp. 225–246, Sep. 1969.
- [18] W. D. Voiers, "Evaluating processed speech using the diagnostic rhyme test," *Speech Technol.*, vol. Jan/Feb, pp. 30–39, 1983.
- [19] "Perceptual evaluation of speech quality (PESQ), and objective method for end-to-end speech quality assessment of narrowband telephone networks and speech codecs," ITU-T Rec., 2000, pp. 862–862.
- [20] M. Cooke and P. Green, "Recognition of speech separated from acoustic mixtures," *J. Acoust. Soc. Amer.*, no. 96, p. 3293, 1994.
- [21] M. Cooke, P. Green, L. Josifovski, and A. Vizinho, "Robust automatic speech recognition with missing and uncertain acoustic data," *Speech Commun.*, vol. 34, pp. 267–285, 2001.

- [22] G. Brown and M. Cooke, "Computational auditory scene analysis," *Comput. Speech Lang.*, vol. 8, pp. 297–336, 1994.
- [23] D. Wang and G. Brown, *Computational Auditory Scene Analysis: Principles, Algorithms, and Applications*. Hoboken, NJ: Wiley, 2006.
- [24] D. Ellis, "Model-based scene analysis," in *Computational Auditory Scene Analysis: Principles, Algorithms, and Applications*, D. Wang and G. Brown, Eds. Hoboken, NJ: Wiley, 2006, pp. 115–146.
- [25] R. Weiss and D. Ellis, "Speech separation using speaker-adapted eigen-voice speech models," *Comput. Speech Lang.*, vol. 24, no. 1, pp. 16–29, 2010.
- [26] D. Wang, "Feature-based speech segregation," in *Computational Auditory Scene Analysis: Principles, Algorithms, and Applications*, D. Wang and G. Brown, Eds. Hoboken, NJ: Wiley, 2006, pp. 81–114.
- [27] D. Wang, "On ideal binary mask as the computational goal of auditory scene analysis," in *Speech Separation by Humans and Machines*, P. Divenyi, Ed. Norwell, MA: Kluwer, 2005, pp. 181–197.
- [28] D. Brungart, P. Chang, B. Simpson, and D. Wang, "Isolating the energetic component of speech-on-speech masking with ideal time–frequency segregation," *J. Acoust. Soc. Amer.*, vol. 120, no. 6, pp. 4007–4018, 2006.
- [29] N. Li and P. Loizou, "Factors influencing intelligibility of ideal binary-masked speech: Implications for noise reduction," *J. Acoust. Soc. Amer.*, vol. 123, no. 3, pp. 1673–1682, 2008.
- [30] D. Wang, U. Kjems, M. S. Pedersen, J. B. Boldt, and T. Lunner, "Speech intelligibility in background noise with ideal binary time–frequency masking," *J. Acoust. Soc. Amer.*, vol. 125, pp. 2336–2347, 2009.
- [31] Y. Hu and P. Loizou, "A new sound coding strategy for suppressing noise in cochlear implants," *J. Acoust. Soc. Amer.*, vol. 124, no. 1, pp. 498–509, 2008.
- [32] S. Araki, S. Makino, H. Sawada, and R. Mukai, "Reducing musical noise by a fine-shift overlap-add method applied to source separation using a time–frequency mask," in *Proc. IEEE Int. Conf. Acoust., Speech, Signal Process.*, 2005, vol. 3, pp. 81–84.
- [33] T. Jan, W. Wang, and D. Wang, "A multistage approach for blind separation of convolutive speech mixtures," in *Proc. IEEE Int. Conf. Acoust., Speech, Signal Process.*, 2009, pp. 1713–1716.
- [34] G. Kim, Y. Lu, Y. Hu, and P. Loizou, "An algorithm that improves speech intelligibility in noise for normal-hearing listeners," *J. Acoust. Soc. Amer.*, vol. 126, no. 3, pp. 1486–1494, 2009.
- [35] Y. Li and D. Wang, "On the optimality of ideal time–frequency masks," in *Speech Commun.*, 2009, vol. 51, pp. 230–239.
- [36] K. Kryter, "Validation of the articulation index," *J. Acoust. Soc. Amer.*, vol. 34, pp. 1698–1706, 1962.
- [37] N. Li and P. Loizou, "Factors influencing intelligibility of ideal binary-masked speech: Implications for noise reduction," *J. Acoust. Soc. Amer.*, vol. 123, no. 3, pp. 1673–1682, 2009.
- [38] Y. Hu and P. Loizou, "Techniques for estimating the ideal binary mask," in *Proc. 11th Int. Workshop Acoust. Echo Noise Control*, 2008.
- [39] A. Drygajlo, "Detection of reliable features for speech recognition in noisy conditions using a statistical criterion," in *Proc. Consistent Rel. Acoust. Cues Sound Anal. Workshop*, 2001, vol. 1, pp. 71–74.
- [40] Y. Ephraim and D. Malah, "Speech enhancement using a minimum mean-square error log-spectral amplitude estimator," *IEEE Trans. Acoust., Speech, Signal Process.*, vol. ASSP-33, no. 2, pp. 443–445, Apr. 1985.
- [41] C. V. Pavlovic, G. A. Studebaker, and R. Sherbecoe, "An articulation index based procedure for predicting the speech recognition performance of hearing-impaired individuals," *J. Acoust. Soc. Amer.*, vol. 80, pp. 50–57, 1986.
- [42] *ANSI S3.5-1997, Methods for Calculation of the Speech Intelligibility Index*, ANSI S3.5-1997, American National Standards Inst., 1997.
- [43] K. Rhebergen, N. Versfeld, and W. Dreschler, "Extended speech intelligibility index for the prediction of the speech reception threshold in fluctuating noise," *J. Acoust. Soc. Amer.*, vol. 120, pp. 3988–3997, 2006.
- [44] K. Kryter, "Methods for calculation and use of the articulation index," *J. Acoust. Soc. Amer.*, vol. 34, no. 11, pp. 1689–1697, 1962.
- [45] N. R. French and J. C. Steinberg, "Factors governing the intelligibility of speech sounds," *J. Acoust. Soc. Amer.*, vol. 19, pp. 90–119, 1947.
- [46] C. V. Pavlovic, "Derivation of primary parameters and procedures for use in speech intelligibility predictions," *J. Acoust. Soc. Amer.*, vol. 82, pp. 413–422, 1987.



Philipos C. Loizou (S'90–M'91–SM'04) received the B.S., M.S., and Ph.D. degrees in electrical engineering from Arizona State University (ASU), Tempe, in 1989, 1991, and 1995, respectively.

From 1995 to 1996, he was a Postdoctoral Fellow in the Department of Speech and Hearing Science, ASU, working on research related to cochlear implants. He was an Assistant Professor at the University of Arkansas, Little Rock, from 1996 to 1999. He is now a Professor and holder of the Cecil and Ida Green Chair in the Department of Electrical Engineering, University of Texas at Dallas, Richardson. His research interests are in the areas of signal processing, speech processing, and cochlear implants. He is the author of the textbook *Speech Enhancement: Theory and Practice* (CRC, 2007) and coauthor of the textbooks *An Interactive Approach to Signals and Systems Laboratory* (National Instruments, 2008) and *Advances in Modern Blind Signal Separation Algorithms: Theory and Applications* (Morgan & Claypool Publishers, 2010). He is currently an Associate Editor for the *International Journal of Audiology*.

Dr. Loizou is a Fellow of the Acoustical Society of America. He is currently an Associate Editor of the IEEE TRANSACTIONS ON BIOMEDICAL ENGINEERING. He was an Associate Editor of the IEEE TRANSACTIONS ON SPEECH AND AUDIO PROCESSING (1999–2002), IEEE SIGNAL PROCESSING LETTERS (2006–2009), and is currently a member of the Speech Technical Committee of the IEEE Signal Processing Society.



Gibak Kim received the B.S. and M.S. degrees in electronics engineering and the Ph.D. degree in electrical engineering from Seoul National University, Seoul, Korea, in 1994, 1996, and 2007, respectively.

From 1996 to 2000, he was with the Machine Intelligence Group, Department of the Information Technology, LG Electronics, Inc., Seoul, Korea. He also worked at Voiceware, Ltd., from 2000 to 2003, as a Senior Research Engineer involved in the development of the automatic speech recognizer. Since 2007, he has been a Research Associate at the University of

Texas at Dallas, Richardson, working on the development of noise-reduction algorithms that can improve speech intelligibility. His general research interests are in speech enhancement, speech recognition, and microphone array signal processing.