



Causal-Anticausal Decomposition of Speech using Complex Cepstrum for Glottal Source Estimation

Thomas Drugman, Baris Bozkurt, T. Dutoit

► To cite this version:

Thomas Drugman, Baris Bozkurt, T. Dutoit. Causal-Anticausal Decomposition of Speech using Complex Cepstrum for Glottal Source Estimation. *Speech Communication*, 2011, 53 (6), pp.855. 10.1016/j.specom.2011.02.004 . hal-00746109

HAL Id: hal-00746109

<https://hal.science/hal-00746109>

Submitted on 27 Oct 2012

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.

Accepted Manuscript

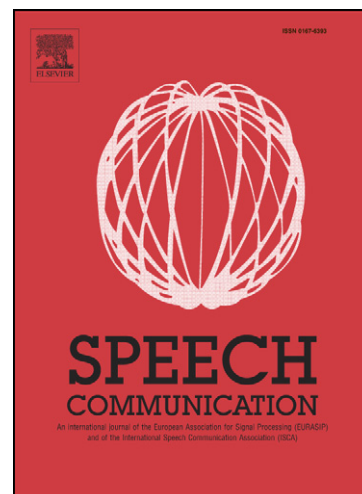
Causal-Anticausal Decomposition of Speech using Complex Cepstrum for Glottal Source Estimation

Thomas Drugman, Baris Bozkurt, Thierry Dutoit

PII: S0167-6393(11)00022-7
DOI: [10.1016/j.specom.2011.02.004](https://doi.org/10.1016/j.specom.2011.02.004)
Reference: SPECOM 1973

To appear in: *Speech Communication*

Received Date: 23 February 2010
Revised Date: 8 February 2011
Accepted Date: 9 February 2011



Please cite this article as: Drugman, T., Bozkurt, B., Dutoit, T., Causal-Anticausal Decomposition of Speech using Complex Cepstrum for Glottal Source Estimation, *Speech Communication* (2011), doi: [10.1016/j.specom.2011.02.004](https://doi.org/10.1016/j.specom.2011.02.004)

This is a PDF file of an unedited manuscript that has been accepted for publication. As a service to our customers we are providing this early version of the manuscript. The manuscript will undergo copyediting, typesetting, and review of the resulting proof before it is published in its final form. Please note that during the production process errors may be discovered which could affect the content, and all legal disclaimers that apply to the journal pertain.

Causal-Anticausal Decomposition of Speech using Complex Cepstrum for Glottal Source Estimation

Thomas Drugman^a, Baris Bozkurt^b, Thierry Dutoit^a

^a*TCTS Lab, University of Mons, Belgium*

^b*Department of Electrical & Electronics Engineering, Izmir Institute of Technology, Turkey*

Abstract

Complex cepstrum is known in the literature for linearly separating causal and anticausal components. Relying on advances achieved by the Zeros of the Z-Transform (ZZT) technique, we here investigate the possibility of using complex cepstrum for glottal flow estimation on a large-scale database. Via a systematic study of the windowing effects on the deconvolution quality, we show that the complex cepstrum causal-anticausal decomposition can be *effectively* used for glottal flow estimation when specific windowing criteria are met. It is also shown that this complex cepstral decomposition gives similar glottal estimates as obtained with the ZZT method. However, as complex cepstrum uses FFT operations instead of requiring the factoring of high-degree polynomials, the method benefits from a much higher speed. Finally in our tests on a large corpus of real expressive speech, we show that the proposed method has the potential to be used for voice quality analysis.

Keywords: Complex Cepstrum, Homomorphic Analysis, Glottal Source Estimation, Source-Tract Separation.

1. Introduction

Glottal source estimation aims at isolating the glottal flow contribution directly from the speech waveform. For this, most of the methods proposed in the literature are based on an inverse filtering process. These methods first estimate a parametric model of the vocal tract, and then obtain the glottal flow by removing the vocal tract contribution via inverse filtering. The methods in this category differ by the way the vocal tract is estimated.

In some approaches [1], [2], this estimation is computed during the glottal closed phase, as the effects of the subglottal cavities are minimized during this period, providing a better way for estimating the vocal tract transfer function. Some other methods (such as [3]) are based on iterative and/or adaptive procedures in order to improve the quality of the glottal flow estimation. Note that a detailed overview of the glottal source estimation methods can be found in various resources such as [4] or [5].

In this paper we consider a non-parametric decomposition of the speech signal based on the mixed-phase model [6],[7]. According to this model, speech contains a maximum-phase (i.e anticausal) component corresponding to the glottal open phase. In a previous work [8], we proposed an algorithm based on the Zeros of the Z-Transform (ZTT) which has the ability to achieve such a deconvolution. However, the ZZT method suffers from high computational load due to the necessity of factorizing large degree polynomials. It has also been discussed in previous studies that the complex cepstrum had the potential to be used for excitation analysis ([9],[10]) but no technique is yet available for reliable glottal flow estimation. This paper more specifically discusses the use of the complex cepstrum for performing the estimation of the glottal open phase from the speech signal, in the light of our previous work on ZZT-based source separation. Almost identical results are obtained with limited computational load, and it is shown that the algorithm is stable enough to enable the analysis of a large database. This manuscript extends our first experiments on such a cepstral decomposition of speech [11] by providing a more comprehensive theoretical framework, by performing extensive tests on a large real speech corpus and by giving access to a freely available Matlab toolbox.

The goal of this paper is two-fold. First we explain in which conditions complex cepstrum can be used for glottal source estimation. The link with the ZZT-based technique is emphasized and both methods are shown to be two means of achieving the same operation: the causal-anticausal decomposition. However it is shown that the complex cepstrum performs it in a much faster way. Secondly the effects of windowing are studied in a systematic framework. This leads to a set of constraints on the window so that the resulting windowed speech segment exhibits properties described by the mixed-phase model of speech. It should be emphasized that no method is here proposed for estimating the return phase component of the glottal flow signal. As the glottal return phase has a causal character [7], its contribution is mixed in the also causal vocal tract filter contribution of the speech signal.

The paper is structured as follows. Section 2 presents the theoretical framework for the causal-anticausal decomposition of voiced speech signals. Two algorithms achieving this deconvolution, namely the Zeros of the Z-Transform (ZZT) and the Complex Cepstrum (CC) based techniques, are described in Section 3. The influence of windowing on the causal-anticausal decomposition is investigated in Section 4 by a systematic study on synthetic signals. Relying on the conclusions of this study, it is shown in Section 5 that the complex cepstrum can be efficiently used for glottal source estimation on real speech. Among others we demonstrate the potential of this method for voice quality analysis on an expressive speech corpus. Finally Section 6 concludes and summarizes the contributions of the paper.

2. Causal-Anticausal Decomposition of Voiced Speech

2.1. Mixed-Phase Model of Voiced Speech

It is generally accepted that voiced speech results from the excitation of a linear time-invariant system with impulse response $h(n)$, by a periodic pulse train $p(n)$ [10]:

$$x(n) = p(n) \star h(n) \quad (1)$$

According to the mechanism of voice production, speech is considered as the result of a glottal flow signal filtered by the vocal tract cavities and radiated by the lips. The system transfer function $H(z)$ then consists of the three following contributions:

$$H(z) = A \cdot G(z)V(z)R(z) \quad (2)$$

where A is the source gain, $G(z)$ the glottal flow over a single cycle, $V(z)$ the vocal tract transmittance and $R(z)$ the radiation load. The resonant vocal tract contribution is generally represented for "pure" vowels by a set of minimum-phase poles ($|v_{2,k}| < 1$), while modeling nasalized sounds requires to also consider minimum-phase (i.e causal) zeros ($|v_{1,k}| < 1$). $V(z)$ can then be written as the rational form:

$$V(z) = \frac{\prod_{k=1}^M (1 - v_{1,k}z^{-1})}{\prod_{k=1}^N (1 - v_{2,k}z^{-1})} \quad (3)$$

During the production of voiced sounds, the airflow evicted by the lungs arises in the trachea and causes a quasi-periodic vibration of the vocal folds

[10]. These latter are then subject to quasi-periodic opening/closure cycles. During the *open phase*, vocal folds are progressively displaced from their initial state because of the increasing subglottal pressure [12]. When the elastic displacement limit is reached, they suddenly return to this position during the so-called *return phase*. Figure 1 displays one cycle of a typical waveform of the glottal flow derivative according to the Liljencrants-Fant (LF) model [13]. The limits of these two phases are indicated on the plot, as well as the particular event separating them, called Glottal Closure Instant (GCI).

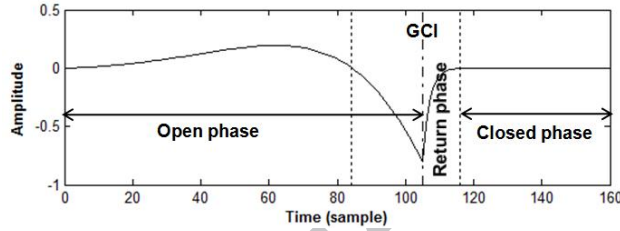


Figure 1: One cycle of a typical waveform of the glottal flow derivative, following the Liljencrants-Fant (LF) model. The different phases of the glottal cycle, as well as the Glottal Closure Instant (GCI) are also indicated.

It has been shown in [14], [7] that the glottal open phase can be modeled by a pair of maximum-phase (i.e anticausal) poles ($|g_2| > 1$) producing the so-called *glottal formant*, while the return phase can be assumed to be a first order causal filter response ($|g_1| < 1$) resulting in a *spectral tilt*:

$$G(z) = \frac{1}{(1 - g_1 z^{-1})(1 - g_2 z^{-1})(1 - g_2^* z^{-1})} \quad (4)$$

As for the lip radiation, its action is generally assumed as a differential operator:

$$R(z) = 1 - r z^{-1} \quad (5)$$

with r close to 1. For this reason, it is generally preferred to consider $G(z)R(z)$ in combination, and consequently to study the *glottal flow derivative* or *differentiated glottal flow* instead of the glottal flow itself.

Gathering the previous equations, the system z -transform $H(z)$ can be expressed as a rational fraction with general form [9]:

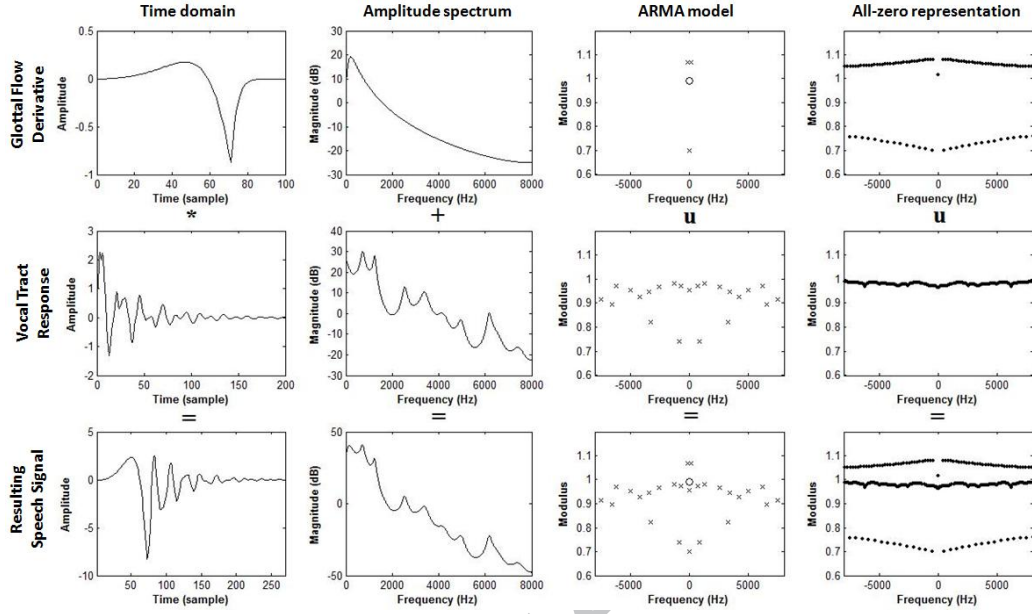


Figure 2: Illustration of the mixed-phase model. The three rows respectively correspond to the glottal flow derivative, the vocal tract response, and the resulting voiced speech. These three signals are all represented in four domains (from the left to the right): waveform, amplitude spectrum, pole-zero modeling, and all-zero (or ZZT) representation. Each column shows how voiced speech is obtained, in each of the four domains.

$$H(z) = A \frac{\prod_{k=1}^{M_i} (1 - a_k z^{-1})}{\prod_{k=1}^{N_i} (1 - b_k z^{-1}) \prod_{k=1}^{N_o} (1 - c_k z^{-1})} \quad (6)$$

where a_k and b_k respectively denote the zeros and poles inside the unit circle ($|a_k|$ and $|b_k| < 1$), while c_k are the poles outside the unit circle ($|c_k| > 1$). The basic idea behind using causal-anticausal decomposition for glottal flow estimation is the following: *since c_k are only related to the glottal flow, isolating the maximum-phase (i.e anticausal) component of voiced speech should then give an estimation of the glottal open phase.* Besides, if the glottal return phase can be considered as abrupt and if the glottal closure is complete, the anticausal contribution of speech corresponds to the glottal flow. If this is not the case [15], these latter components are causal (given their damped nature) and the anticausal contribution of voiced speech still gives an estimation of the glottal open phase.

Figure 2 illustrates the mixed-phase model on a single frame of synthetic

vowel. In each row the glottal flow and vocal tract contributions, as well as the resulting speech signal, are shown in a different representation space. It should be emphasized here that the all-zero representation (later referred to as the Zeros of Z-Transform (ZZT) representation, and shown in the last column) is obtained by a root finding operation (i.e. a finite(n)-length signal frame is represented with only zeros in the z -domain). There exists $n - 1$ zeros (of the z -transform) for a signal frame with n samples. However the zero in the third row comes from the ARMA model and hence should not be confused with the ZZT. The first row shows a typical glottal flow derivative signal. From the ZZT representation (last column), it can be noticed that some zeros lie outside the unit circle while others are located inside it. The outside zeros correspond to the maximum-phase glottal opening, while the others come from the minimum-phase glottal closure [8]. The vocal tract response is displayed in the second row. All its zeros are inside the unit circle due to its damped exponential character. Finally the last row is related to the resulting voiced speech. Interestingly its set of zeros is simply the union of the zeros of the two previous components. This is due to the fact that the convolution operation in the time domain corresponds to the multiplication of the z -transform polynomials in the z -domain. For a detailed study of ZZT representation and the mixed-phase speech model, the reader is referred to [8].

2.2. Short-Time Analysis of Voiced Speech

For real speech data, Equation (1) is only valid for a short-time signal [16], [17]. Most practical applications therefore require processing of windowed (i.e short-time) speech segments:

$$s(n) = w(n)x(n) \quad (7)$$

$$= w(n)(A \cdot p(n) \star g(n) \star v(n) \star r(n)) \quad (8)$$

and the goal of the decomposition is to extract the glottal source component $g(n)$ from $s(n)$. As it will be discussed throughout this article, windowing is of crucial importance in order to achieve a correct deconvolution. Indeed, the z -transform of $s(n)$ can be written as:

$$S(z) = W(z) \star X(z) \quad (9)$$

$$= \sum_{n=0}^{N-1} w(n)x(n)z^{-n} \quad (10)$$

$$= s(0)z^{-N+1} \prod_{k=1}^{M_i} (z - Z_{C,k}) \prod_{k=1}^{M_o} (z - Z_{AC,k}) \quad (11)$$

where Z_C and Z_{AC} are respectively a set of M_i causal ($|Z_{C,k}| < 1$) and M_o anticausal ($|Z_{AC,k}| > 1$) zeros (with $M_o + M_i = N - 1$). As it will be underlined in Section 3.1, Equation 11 corresponds to the ZZT representation.

From these latter expressions, two important considerations have now to be taken into account:

- Since $s(n)$ is finite length, $S(z)$ is a polynomial in z (see Eq. 11). This means that the poles of $H(z)$ are now embedded under an all-zero form. Indeed let us consider a single real pole a . The z-transform of the related impulse response $y(n)$ limited to N points is [18]:

$$Y(z) = \sum_{n=0}^{N-1} a^n z^{-n} = \frac{1 - (az^{-1})^N}{1 - az^{-1}} \quad (12)$$

which is an all-zero form, since the root of the denominator is also a root of the numerator (and the pole is consequently cancelled).

- It can be seen from Equations (9) and (10) that the window $w(n)$ may have a dramatic influence on $S(z)$ [17], [10]. As windowing in the time domain results in a convolution of the window spectrum with the speech spectrum, the resulting change in the ZZT is a highly complex issue to study [19]. Indeed the multiplication by the windowing function (as in Equation 10) modifies the root distribution of $X(z)$ in a complex way that cannot be studied analytically. For this reason, the impact of the windowing effects on the mixed-phase model is studied in this paper in an empirical way, as it was done in [17] and [10] for the convolutional model.

To emphasize the crucial role of windowing, Figures 3 and 4 respectively display a case of correct and erroneous glottal flow estimation via causal-anticausal decomposition on a real speech segment. In these figures, the top-left panel (a) contains the speech signal together with the applied window and the synchronized differenced ElectroGlottograph *dEGG* (after compensation of the delay between the laryngograph and the microphone). Peaks in the *dEGG* signal are informative about the location of the Glottal Closure Instant (GCI). The top-right panel (b) plots the roots of the windowed signal ($Z_{C,k}$ and $Z_{AC,k}$) in polar coordinates. The bottom panels (c) and (d) correspond to the time waveform and amplitude spectrum of the maximum-phase (i.e. anticausal) component which is expected to correspond to the glottal flow open phase.

In Figure 3, an appropriate window respecting the conditions we will derive in Section 4 is used. This results in a good separation between the zeros inside and outside the unit circle (see Fig.3(b)). The windowed signal then exhibits good mixed-phase properties and the resulting maximum and minimum-phase components corroborate the model exposed in Section 2.1. On the contrary, a 25 ms long Hanning window is employed in Figure 4, as widely used in speech processing. It can be seen that even when this window is centered on a GCI, the resulting causal-anticausal decomposition is erroneous. Zeros on each side of the unit circle are not well separated: the windowed signal does not exhibit characteristics of the mixed-phase model. This simple comparison highlights the dramatic influence of windowing on the deconvolution. In Section 4, we discuss in detail the set of properties the window should convey so as to yield a good decomposition.

3. Algorithms for Causal-Anticausal Decomposition of Voiced Speech

For a segment $s(n)$ resulting from an appropriate windowing of a voiced speech signal $x(n)$, two algorithms are compared for achieving causal-anticausal decomposition, thereby leading to an estimate $\tilde{g}(n)$ of the real glottal source $g(n)$. The first one relies on the Zeros of the Z-Transform (ZZT, [8]) and is summarized in Section 3.1. The second technique is based on the Complex Cepstrum (CC) and is described in Section 3.2. It is important to note that both methods are functionally equivalent to each other, in the sense that they take the same input $s(n)$ and should give the same output $\tilde{g}(n)$. As emphasized in Section 2.2, the quality of the decomposition then only depends on the applied windowing, i.e. whether $s(n) = w(n)x(n)$ exhibits expected

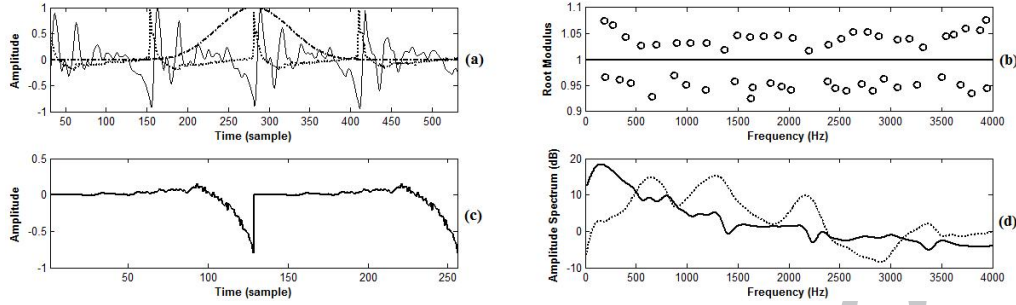


Figure 3: Example of decomposition on a real speech segment using an appropriate window. (a): The speech signal (solid line) with the synchronized dEGG (dotted line) and the applied window (dash-dotted line). (b): The zero distribution in polar coordinates. (c): Two cycles of the maximum-phase component (corresponding to the glottal flow open phase). (d): Amplitude spectra of the minimum (dotted line) and maximum-phase (solid line) components of the speech signal. It can be observed that the windowed signal respects the mixed-phase model since the zeros on each side of the unit circle are well separated.

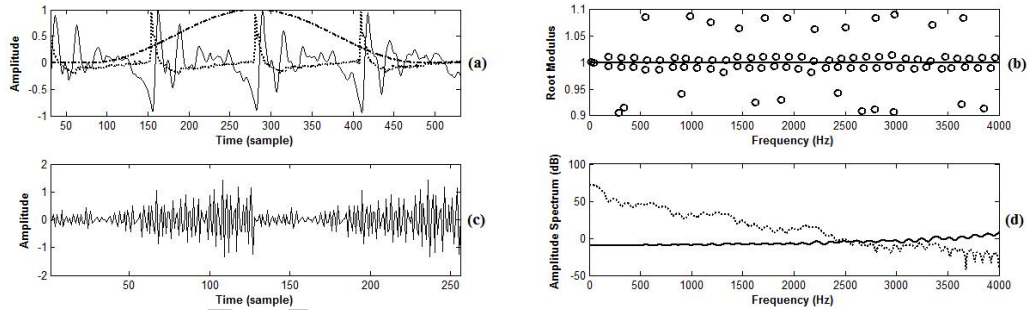


Figure 4: Example of decomposition on a real speech segment using a 25 ms long Hanning window. (a): The speech signal (solid line) with the synchronized dEGG (dotted line) and the applied window (dash-dotted line). (b): The zero distribution in polar coordinates. (c): Two cycles of the maximum-phase component. (d): Amplitude spectra of the minimum (dotted line) and maximum-phase (solid line) components of the speech signal. The zeros on each side of the unit circle are not well separated and the windowed signal does not respect the mixed-phase model. The resulting deconvolved components are irrelevant (while their convolution still gives the input speech signal).

mixed-phase properties or not. It will then be shown that both methods lead to similar results (see Section 5.2). However, on a practical point of view, the use of the complex cepstrum is advantageous since it will be shown that it is much faster than ZZT. Note that we made a Matlab toolbox containing

these two methods freely available in <http://tcts.fpm.s.ac.be/~drugman/>.

3.1. Zeros of the Z-Transform-based Decomposition

According to Equation (11), $S(z)$ is a polynomial in z with zeros inside and outside the unit circle. The idea of the ZZT-based decomposition is to isolate the roots Z_{AC} and to reconstruct from them the anticausal component. The algorithm can then be summarized as follows [8]:

1. *Window the signal with guidelines provided in Section 4,*
2. *Compute the roots of the polynomial $S(z)$,*
3. *Isolate the roots with a modulus greater than 1,*
4. *Compute $\tilde{G}(z)$ from these roots.*

Although very simple, this technique requires the factorization of a polynomial whose order is generally high (depending on the sampling rate and window length). Even though current factoring algorithms are accurate, the time complexity still remains high [20].

In addition to [8] where the ZZT algorithm is introduced, some recent studies [21], [22] have shown that ZZT outperforms other well-known methods of glottal flow estimation in clean recordings. Its main disadvantages are reported as sensitivity to noise and high computational load.

3.2. Complex Cepstrum-based Decomposition

Homomorphic systems have been developed in order to separate non-linearly combined signals [9]. As a particular example, the case where inputs are convolved is especially important in speech processing. Separation can then be achieved by a linear homomorphic filtering in the complex cepstrum domain, which interestingly presents the property to map time-domain convolution into addition. In speech analysis, complex cepstrum is usually employed to deconvolve the speech signal into a periodic pulse train and the vocal system impulse response [10], [17]. It finds applications such as pitch detection [23], vocoding [24], etc. Based on a previous study [11], it is here detailed how to use the complex cepstrum in order to estimate the glottal flow by achieving the causal-anticausal decomposition introduced in Section 2.2. To our knowledge, no complex cepstrum-based glottal flow estimation method is available in the literature (except this manuscript's introductory version [11]). Hence it is one of the novel contributions of this paper to introduce one and to test it on a large real speech database.

The complex cepstrum (CC) $\hat{s}(n)$ of a discrete signal $s(n)$ is defined by the following equations [9]:

$$S(\omega) = \sum_{n=-\infty}^{\infty} s(n)e^{-j\omega n} \quad (13)$$

$$\log[S(\omega)] = \log(|S(\omega)|) + j\angle S(\omega) \quad (14)$$

$$\hat{s}(n) = \frac{1}{2\pi} \int_{-\pi}^{\pi} \log[S(\omega)] e^{j\omega n} d\omega \quad (15)$$

where Equations (13), (14) and (15) are respectively the Discrete-Time Fourier Transform (DTFT), the complex logarithm and the inverse DTFT (IDTFT). One difficulty when computing the CC lies in the estimation of $\angle S(\omega)$, which requires an efficient phase unwrapping algorithm. In this work, we computed the FFT on a sufficiently large number of points (typically 4096) such that the grid on the unit circle is sufficiently fine to facilitate in this way the phase evaluation.

If $S(z)$ is written as in Equation (11), it can be easily shown [9] that the corresponding complex cepstrum can be expressed as:

$$\hat{s}(n) = \begin{cases} |s(0)| & \text{for } n = 0 \\ \sum_{k=1}^{M_o} \frac{Z_{AC,k}^n}{n} & \text{for } n < 0 \\ \sum_{k=1}^{M_i} \frac{Z_{C,k}^n}{n} & \text{for } n > 0 \end{cases} \quad (16)$$

This equation shows the close link between the ZZT and the CC-based techniques. Relying on this equation, Steiglitz and Dickinson demonstrated the possibility of computing the complex cepstrum and unwrapped phase by factoring the z-transform [25], [26]. The approach we propose is just the inverse thought process in the sense that our goal is precisely to use the complex cepstrum in order to avoid any factorization. In this way we show that the complex cepstrum can be used as an efficient means to estimate the glottal flow, while circumventing the requirement of factoring polynomials (as it is the case for the ZZT). Indeed it will be shown in Section 4.2 that optimal windows have their length proportional to the pitch period. The ZZT-based technique then requires to compute the roots of generally high-order polynomials (depending on the sampling rate and on the pitch). Although current polynomial factoring algorithms are accurate, the computational load still remains high, with a complexity order of $O(n^2)$ for the fastest algorithms

[20], where n denotes the number of samples in the considered frame. On the other hand, the CC-based method just relies on FFT and IFFT operations which can be fast computed, and whose order is $O(N_{FFT} \log(N_{FFT}))$, where N_{FFT} is fixed to 4096 in this work for facilitating phase unwrapping, as mentioned above. For this reason a change in the frame length has little influence on the computation time for the CC-based method. Table 1 compares both methods in terms of computation time. The use of the complex cepstrum now offers the possibility of integrating a causal-anticausal decomposition module into a real-time application, which was previously almost impossible with the ZZT-based technique.

Pitch	ZZT-based decomposition	CC-based decomposition
60 Hz	111.4	1.038
180 Hz	11.2	1

Table 1: Comparison of the relative computation time (for our Matlab implementation with $F_s = 16kHz$) required for decomposing a two pitch period long speech frame. Durations were normalized according to the time needed by the complex cepstrum-based deconvolution for $F_0 = 180Hz$.

Regarding Equation (16), it is obvious that causal-anticausal decomposition can be performed using the complex cepstrum, as follows [11]:

1. *Window the signal with guidelines provided in Section 4,*
2. *Compute the complex cepstrum $\hat{s}(n)$ using Equations (13), (14) and (15),*
3. *Set $\hat{s}(n)$ to zero for $n > 0$,*
4. *Compute $\tilde{g}(n)$ by applying the inverse operations of Equations (13), (14) and (15) on the resulting complex cepstrum.*

Figure 5 illustrates the complex cepstrum-based decomposition for the example shown in Figure 3. A simple linear liftering keeping only the negative (positive) indexes of the complex cepstrum allows to isolate the maximum and minimum phase components of voiced speech. It should be emphasized that windowing is very critical as it is the case for the ZZT decomposition. The example in Figure 4 (where a 25ms long Hanning window is used) would lead to an unsuccessful decomposition. We think that this critical dependence on the window function, length and location was the main hindrance in developing a complex cepstrum-based glottal flow estimation method, although the potential is known earlier in the literature [10].

It is also worth noting that since the CC method is an alternative means of achieving the mixed-phase decomposition, it suffers from the same noise sensitivity as the ZZT does.

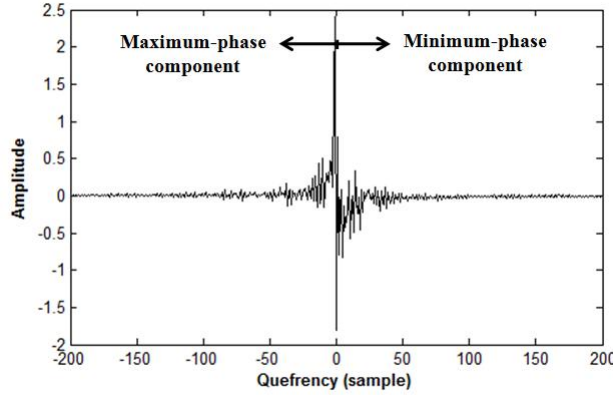


Figure 5: The complex cepstrum $\hat{s}(n)$ of the windowed speech segment $s(n)$ presented in Figure 3(a). The maximum- (minimum-) phase component can be isolated by only considering the negative (positive) indexes of the complex cepstrum.

4. Experiments on Synthetic Speech

The goal of this section is to study, on synthetic speech signals, the impact of the windowing effects on the causal-anticausal decomposition. It is one of the main contributions of this study to provide a parametric analysis of the windowing problem and provide guidelines for reliable complex cepstrum-based glottal flow estimation. For this, synthetic speech signals are generated for a wide range of test conditions [11]. The idea is to cover the diversity of configurations one could find in natural speech by varying all parameters over their whole range. Synthetic speech is produced according to the source-filter model by passing a synthetic train of Liljencrants-Fant (LF) glottal waves [13] through an auto-regressive filter extracted by LPC analysis of real sustained vowels uttered by a male speaker. As the mean pitch in these utterances is about 100 Hz, it is reasonable to consider that the fundamental frequency should not exceed 60 and 180 Hz in continuous speech. Experiments in this section can then be seen as a proof of concept on synthetic male speech. Table 2 summarizes all test conditions.

Decomposition quality is assessed through two objective measures [11]:

Pitch	60:20:180 Hz
Open quotient	0.4:0.05:0.9
Asymmetry coefficient	0.6:0.05:0.9
Vowel	/a/, /@/, /i/, /y/

Table 2: Table of synthesis parameter variation range.

- **Spectral distortion** : Many frequency-domain measures for quantifying the distance between two speech frames have been proposed in the speech coding literature [27]. A simple relevant measure between the estimated $\hat{g}(n)$ and the real glottal pulse $g(n)$ is the spectral distortion (SD) defined as [27]:

$$SD(g, \hat{g}) = \sqrt{\int_{-\pi}^{\pi} (20 \log_{10} |\frac{G(\omega)}{\hat{G}(\omega)}|)^2 \frac{d\omega}{2\pi}} \quad (17)$$

where $G(\omega)$ and $\hat{G}(\omega)$ denote the DTFT of the original target glottal pulse $g(n)$ and of the estimate $\hat{g}(n)$. To give an idea, it is argued in [28] that a difference of about 1dB (with a sampling rate of 8kHz) is rather imperceptible.

- **Glottal formant determination rate** : The amplitude spectrum for a voiced source generally presents a resonance called the *glottal formant* ([31], see also Section 2.1). As this parameter is an essential feature of the glottal open phase, an error on its determination after decomposition should be penalized. For this, we define the *glottal formant determination rate* as the proportion of frames for which the relative error on the glottal formant frequency is lower than 10%.

This formal experimental protocol allows us to reliably assess our technique and to test its sensivity to various factors influencing the decomposition, such as the window location, function and length. Indeed, Tribolet et al. already observed in 1977 that the window shape and onset may lead to zeros whose topology can be detrimental for accurate pulse estimation [16]. The goal of this empirical study on synthetic signals is precisely to handle with these zeros close to the unit circle, such that the applied window leads to a correct causal-anticausal separation.

4.1. Influence of the window location

In [10] the need of aligning the window center with the system response is highlighted. Analysis is then performed on windows centered on GCIs, as these particular events demarcate the boundary between the causal and anticausal responses, and the linear phase contribution is removed. Figure 6 illustrates the sensitivity of the causal-anticausal decomposition to the window position. It can be noticed that the performance rapidly degrades, especially if the window is centered on the left of the GCI. It is then recommended to apply a GCI-centered windowing. In a concrete application, techniques like the DYPSA algorithm [29], or the method we proposed in [30], have been shown to give a reliable and accurate estimation of the GCI locations directly from the speech signal. For cases for which GCI information is not available or unreliable, the formalism of the mixed-phase separation has been extended in [44] to a chirp analysis, allowing the deconvolution to be achieved in an asynchronous way, but at the expense of a slight performance degradation.

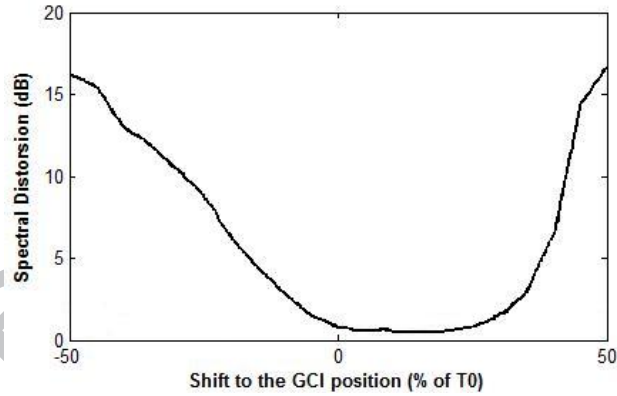


Figure 6: Sensitivity of the causal-anticausal decomposition to a GCI location error. The spectral distortion dramatically increases if a non GCI-centered windowing is applied (particularly on the left of the GCI).

4.2. Influence of the window function and length

In Section 2.2, Figures 3 and 4 showed an example of correct and erroneous decomposition respectively. The only difference between these figures was the length and shape of the applied windowing. To study this effect let

us consider a particular family of windows $w(n)$ of N points satisfying the form [9]:

$$w(n) = \frac{\alpha}{2} - \frac{1}{2} \cos\left(\frac{2\pi n}{N-1}\right) + \frac{1-\alpha}{2} \cos\left(\frac{4\pi n}{N-1}\right) \quad (18)$$

where α is a parameter comprised between 0.7 and 1 (for α below 0.7, the window includes negative values which should be avoided). The widely used Hanning and Blackman windows are particular cases of this family for $\alpha = 1$ and $\alpha = 0.84$ respectively. Figure 7 displays the evolution of the decomposition quality when α and the window length vary. It turns out that a good deconvolution can be achieved as long as the window length is adapted to its shape (or vice versa). For example, the optimal length is about $1.5 T_0$ for a Hanning window and $1.75 T_0$ for a Blackman window. A similar observation can be drawn from Figure 8 according to the spectral distortion criterion. Note that we displayed the inverse spectral distortion $1/SD$ instead of SD only for better viewing purposes. At this point it is interesting to notice that these constraints on the window aiming at respecting the mixed-phase model are sensibly different from those imposed to respect the so-called *convolutional model* [17], [10]. For this latter case, it was indeed recommended to use windows such as Hanning or Hamming with a duration of about 2 to 3 pitch periods. It can be seen from Figure 7 that this would lead to poor causal-anticausal decomposition results. Finally note that it was proposed in [32] to analytically derive the optimal frame length for the causal-anticausal decomposition, by satisfying an immiscibility criterion based on a Cauchy bound.

5. Experiments on Real Speech

The goal of this section is to show that a reliable glottal flow estimation is possible on real speech using the complex cesptrum. The efficiency of this method will be confirmed in Sections 5.1 and 5.2 by analyzing short segments of real speech. Besides we demonstrate in Section 5.3 the potential of using complex cepstrum for voice quality analysis on a large expressive speech corpus.

For these experiments, speech signals sampled at $16kHz$ are considered. The pitch contours are extracted using the Snack library [33] and the glottal closure instants are located directly from the speech waveforms using the algorithm we proposed in [30]. Speech frames are then obtained by applying

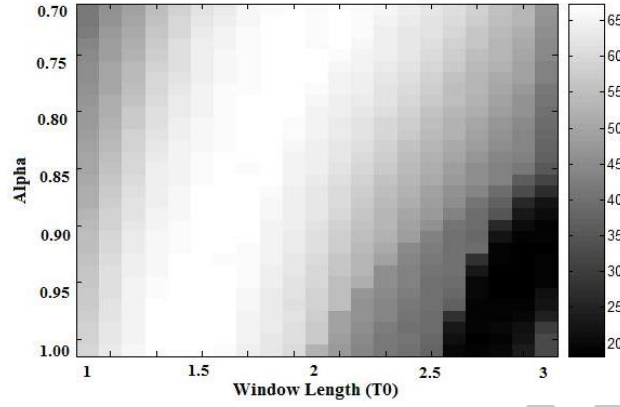


Figure 7: Evolution of the glottal formant determination rate according the window length and shape. Note that the Hanning and Blackman windows respectively correspond to $\alpha = 1$ and $\alpha = 0.84$.

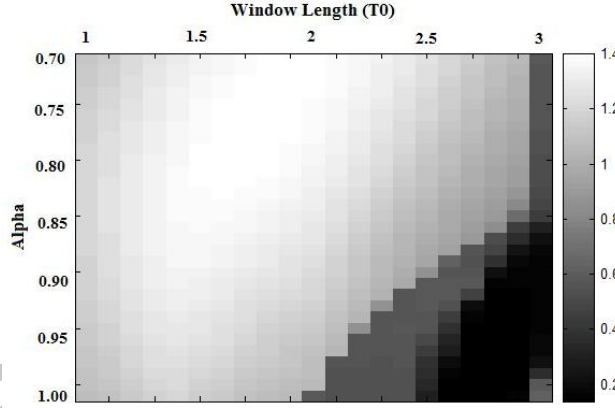


Figure 8: Evolution of the inverse spectral distortion $1/SD$ according the window length and shape. Note that the Hanning and Blackman windows respectively correspond to $\alpha = 1$ and $\alpha = 0.84$. The inverse SD is plotted instead of the SD itself only for clarity purpose.

a GCI-centered windowing. The window we use satisfies Equation (18) for $\alpha = 0.7$ and is two pitch period-long so as to respect the conditions derived in Section 4. Causal-anticausal decomposition is then achieved by the complex cepstrum-based method.

5.1. Example of Decomposition

Figure 9 illustrates a concrete case of decomposition on a voiced speech segment (diphone /*am*/) uttered by a female speaker. It can be seen that even on a nasalized phoneme the glottal source estimation seems to be correctly carried out for most speech frames (i.e the obtained waveforms turn out to corroborate the model of the glottal pulse described in Section 2.1). For some rare cases the causal-anticausal decomposition is erroneous and the maximum-phase component contains a high-frequency irrelevant noise. Nevertheless the spectrum of this maximum-phase contribution almost always presents a low-frequency resonance due to the glottal formant.

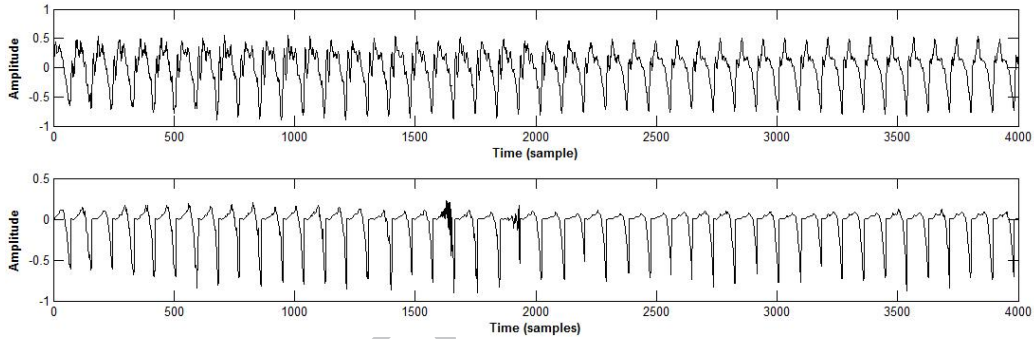


Figure 9: *Top panel:* A segment of voiced speech (diphone /*am*/) uttered by a female speaker. *Bottom panel:* Its corresponding glottal source estimation obtained using the complex cepstrum-based decomposition. It turns out that a reliable estimation can be achieved for most of the speech frames.

5.2. Analysis of sustained vowels

In this experiment, we considered a sustained vowel /*a*/ with a flat pitch which was voluntarily produced with an increasing pressed vocal effort. Here the aim is to show that voice quality variation is reflected as expected on the glottal flow estimates obtained using the causal-anticausal decomposition. Figure 10 plots the evolution of the glottal formant frequency Fg and bandwidth Bw during the phonation [11]. These features were estimated with both ZZT and CC-based methods. It can be observed that, as expected, these techniques lead to similar results. The very slight differences may be due to the fact that, for the complex cepstrum, Equation (16) is realized on a finite number n of points. Another possible explanation is the precision problem in root computation for the ZZT-based technique. In any case, it can

be noticed that the increasing vocal effort can be characterized by increasing values of F_g and B_w .

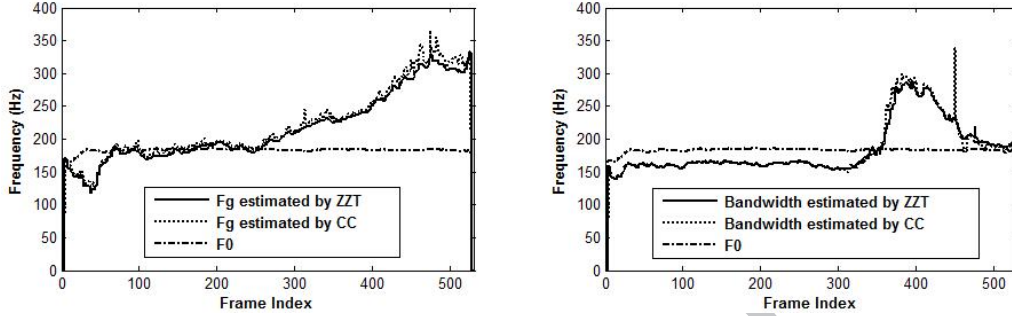


Figure 10: Glottal formant characteristics estimated by both ZZT and CC-based techniques on a real sustained vowel with an increasing pressed effort [11]. *Left panel:* Evolution of the glottal formant frequency. *Right panel:* Evolution of the glottal formant 3dB bandwidth.

5.3. Analysis of an Expressive Speech Corpus

The goal of this part is to show that the differences present in the glottal source when a speaker produces various voice qualities can be tracked using causal-anticausal decomposition. For this, the De7 database is used. This database was designed by Marc Schroeder as one of the first attempts of creating diphone databases for expressive speech synthesis [34]. The database contains three voice qualities (modal, soft and loud) uttered by a German female speaker, with about 50 minutes of speech available for each voice quality.

For each voiced speech frame, the complex cepstrum-based decomposition is performed. The resulting maximum-phase component is then down-sampled at 8kHz and is assumed to give an estimation of the glottal flow derivative for the considered frame. For each segment of voiced speech, a signal similar to the one illustrated in Figure 9 is consequently obtained. For this latter example it was observed that an erroneous decomposition might appear for some frames, leading to an irrelevant high-frequency noise in the estimated anticausal contribution (also observed in Figure 4). One first thing one could wonder is how large is the proportion of such frames over the whole database. As a criterion deciding whether a frame is considered as correctly decomposed or not, we inspect the spectral center of gravity. The distribution of this feature is displayed in Figure 11 for the loud voice. A principal

mode at around 2kHz clearly emerges and corresponds to the majority of frames for which a correct decomposition is carried out. A second minor mode at higher frequencies is also observed. It is related to the frames where the causal-anticausal decomposition fails, leading to a maximum-phase signal containing an irrelevant high-frequency noise (as explained above). It can be noticed from this histogram (and it was confirmed by a manual verification of numerous frames) that fixing a threshold at around 2750 Hz makes a good distinction between frames that are correctly and incorrectly decomposed. According to this criterion, Table 3 summarizes for the whole database the percentage of frames leading to a correct estimation of the glottal flow.

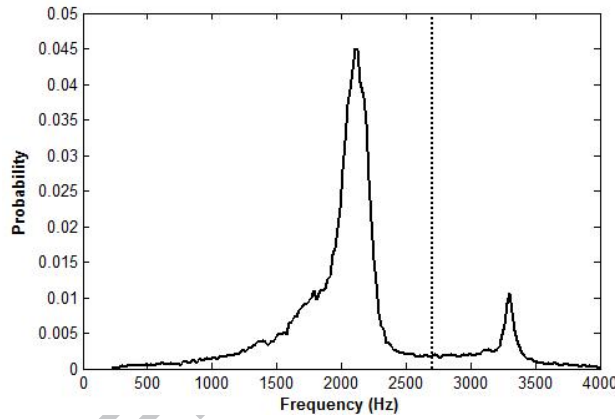


Figure 11: Distribution of the spectral center of gravity of the maximum-phase component, computed for the whole dataset of loud samples. Fixing a threshold around 2.7kHz makes a good separation between correctly and incorrectly decomposed frames.

Voice Quality	% of frames correctly decomposed
Loud	87.22%
Modal	84.41%
Soft	83.69%

Table 3: Proportion of frames leading to a correct causal-anticausal decomposition for the three voice qualities.

For each frame correctly deconvolved, the glottal flow is then characterized by the 3 following common features:

- the Normalized Amplitude Quotient(NAQ): NAQ is a parameter characterizing the glottal closing phase [35]. It is defined as the ratio be-

tween the maximum of the glottal flow and the minimum of its derivative, and then normalized with respect of the fundamental frequency. Its robustness and efficiency to separate different types of phonation was shown in [35]. Note that a quasi-similar feature called *basic shape parameter* was proposed by Fant in [36], where it was qualified as “*most effective single measure for describing voice qualities*”.

- the $H1 - H2$ ratio: This parameter is defined as the ratio between the amplitudes of the amplitude spectrum of the glottal source at the fundamental frequency and at the second harmonic [37], [38]. It has been widely used as a measure characterizing voice quality [39], [36], [4].
- the Harmonic Richness Factor (HRF): This parameter quantifies the amount of harmonics in the magnitude spectrum of the glottal source. It is defined as the ratio between the sum of the amplitudes of harmonics, and the amplitude at the fundamental frequency [12]. It was shown to be informative about the phonation type in [40] and [4].

Figure 12 shows the histograms of these 3 parameters for the three voice qualities. Significant differences between the distributions are observed. Among others it turns out that the production of a louder (softer) voice results in lower (higher) NAQ and $H1 - H2$ values, and of a higher (lower) Harmonic Richness Factor (HRF). These conclusions corroborate the results recently obtained on sustained vowels by Alku in [4] and [35]. Another observation that can be drawn from the histogram of $H1 - H2$ is the presence of two modes for the modal and loud voices. This may be explained by the fact that the estimated glottal source sometimes comprises a ripple both in the time and frequency domains [41]. Indeed consider Figure 13 where two typical cycles of the glottal source are presented for both the soft and loud voice. Two conclusions can be drawn from it. First of all, it is clearly seen that the glottal open phase response for the soft voice is slower than for the loud voice. As it was underlined in the experiment of Section 5.2, this confirms the fact Fg/F_0 increases with the vocal effort. Secondly the presence of a ripple in the loud glottal waveform is highlighted. This has two possible origins: an incomplete separation between Fg and the first formant F_1 [42], and/or a non-linear interaction between the vocal tract and the glottis [41], [43]. This ripple affects the low-frequency contents of the glottal source spectrum, and may consequently perturb the estimation of the $H1 - H2$ feature.

This may therefore explain the second mode in the $H1 - H2$ histogram for the modal and loud voices (where ripple was observed).

It is also observed that histograms in Figure 12 present some overlaps. These overlaps may be explained by the three following reasons. *i)* As histograms result from a study led on a large database of connected speech, the glottal production cannot be expected to be perfectly different as a function of the produced voice quality. *ii)* The parametrization of the glottal waveforms by a single feature can only capture a proportion of their differences. *iii)* It might happen for some speech frames that the glottal estimation fails. Although it is impossible to discern and quantify how much each of these causes explains overlaps in Fig. 12, we believe that the first two reasons are predominant since irrelevant decompositions have been removed using the spectral criterion.

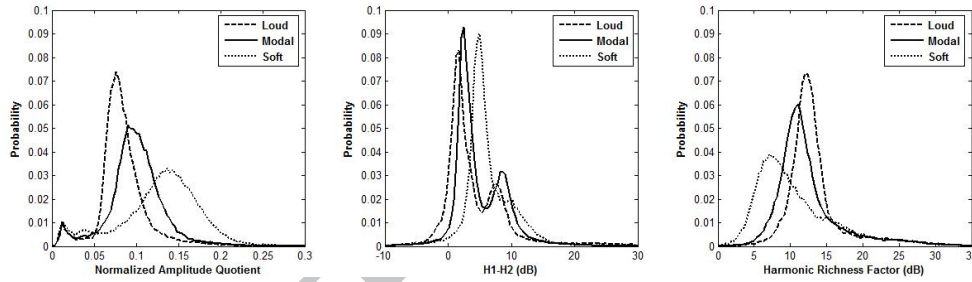


Figure 12: Distributions, computed on a large expressive speech corpus, of glottal source parameters for three voice qualities: (*left*) the Normalized Amplitude Quotient (NAQ), (*middle*) the $H1 - H2$ ratio, and (*right*) the Harmonic Richness Factor (HRF).

6. Discussion and Conclusion

This paper explained the causal-anticausal decomposition principles in order to estimate the glottal source directly from the speech waveform. We showed that the complex cepstrum can be effectively used for this purpose as an alternative to the Zeros of the Z-Transform (ZZT) algorithm. Both techniques were shown to be functionally equivalent to each other, while the complex cepstrum is advantageous for its much higher speed, making it suitable for real-time applications. Windowing effects were studied in a systematic way on synthetic signals. It was emphasized that windowing plays a crucial role. More particularly we derived a set of constraints the window should respect so that the windowed signal matches the mixed-phase model.

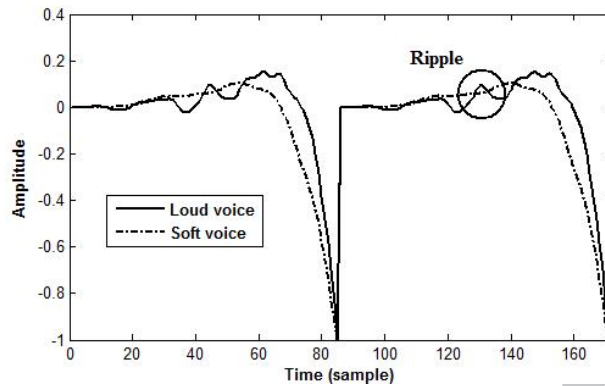


Figure 13: Comparison between two cycles of typical glottal source for both soft (dash-dotted line) and loud voice (solid line). The presence of a ripple in the loud excitation can be observed.

Finally, results on a real speech database (logatoms recorded for the design of an unlimited domain expressive speech synthesizer) were presented for voice quality analysis. The glottal flow was estimated on a large database containing various voice qualities. Interestingly some significant differences between the voice qualities were observed in the excitation. The methods proposed in this paper may be used in several potential applications of speech processing such as emotion detection, speaker recognition, expressive speech synthesis, automatic voice pathology detection and various other applications where real-time glottal source estimation may be useful. Finally note that a Matlab toolbox containing these algorithms is freely available in [http : //tcts.fpms.ac.be/~drugman/](http://tcts.fpms.ac.be/~drugman/).

Acknowledgment

Thomas Drugman is supported by the Fonds National de la Recherche Scientifique (FNRS). Baris Bozkurt is supported by the Scientific and Technological Research Council of Turkey (TUBITAK). The authors also would like to thank N. Henrich and B. Doval for providing us the speech recording used to create Figure 10 and M. Schroeder for the De7 database [34] used in the second experiment on real speech. Authors also would like to thank reviewers for their fruitful feedback.

- [1] D. Veeneman, S. BeMent, *Automatic glottal inverse filtering from speech*

- and electroglottographic signals, IEEE Trans. on Signal Processing, vol. 33, pp. 369-377, 1985.
- [2] P. Alku, E. Vilkman, *Estimation of the glottal pulseform based on discrete all-pole modeling*, Third International Conference on Spoken Language Processing, pp. 1619-1622, 1994.
 - [3] P. Alku, J. Svec, E. Vilkman, F. Sram, *Glottal wave analysis with pitch synchronous iterative adaptive inverse filtering*, Speech Communication, vol. 11, issue 2-3, pp. 109-118, 1992.
 - [4] P. Alku, C. Magi, S. Yrttiaho, T. Bäckström, B. Story, *Closed phase covariance analysis based on constrained linear prediction for glottal inverse filtering*, The Journal of the Acoustical Society of America, vol. 125, issue 5, p. 3289-3305, 2009.
 - [5] J. Walker, P. Murphy, *A Review of Glottal Waveform Analysis*, Progress in Nonlinear Speech Processing, pp. 1-21, 2007.
 - [6] B. Bozkurt, T. Dutoit, *Mixed-phase speech modeling and formant estimation, using differential phase spectrums*, VOQUAL'03, pp. 21-24, 2003.
 - [7] B. Doval, C. D'Alessandro, N. Henrich, *The voice source as a causal/anticausal linear filter*, Proceedings ISCA ITRW VOQUAL03, pp. 15-19, 2003.
 - [8] B. Bozkurt, B. Doval, C. D'Alessandro, T. Dutoit, *Zeros of Z-Transform Representation With Application to Source-Filter Separation in Speech* IEEE Signal Processing Letters, vol. 12, no. 4, 2005.
 - [9] A. Oppenheim, R. Schaffer, *Discrete-time signal processing*, Prentice-Hall, chap. 12, 1989.
 - [10] T. Quatieri, *Discrete-time speech signal processing*, Prentice-Hall, chap. 6, 2002.
 - [11] T. Drugman, B. Bozkurt, T. Dutoit, *Complex Cepstrum-based Decomposition of Speech for Glottal Source Estimation*, Proc. Interspeech, 2009.
 - [12] D. Childers, *Speech Processing and Synthesis Toolboxes*, Wiley and Sons, Inc., 1999.

- [13] G. Fant, J. Liljencrants, Q. Lin, *A four parameter model of glottal flow*, STL-QPSR4, pp. 1-13, 1985.
- [14] W. Gardner, B. Rao, *Noncausal all-pole modeling of voiced speech*, IEEE Trans. on Audio and Speech Processing, vol. 5, issue 1, pp. 1-10, 1997.
- [15] H. Deng, R. Ward, M. Beddoes, M. Hodgson, *A new method for obtaining accurate estimates of vocal-tract filters and glottal waves from vowel sounds*, IEEE Trans. ASSP, vol. 14, pp. 445-455, 2006.
- [16] J. Tribolet, T. Quatieri, A. Oppenheim, *Short-time homomorphic analysis*, Proc. ICASSP77, vol. 2, pp. 716-722, 1977.
- [17] W. Verhelst, O. Steenhaut, *A new model for the short-time complex cepstrum of voiced speech*, IEEE Trans. ASSP, vol. 34, pp. 43-51, 1986.
- [18] A. Oppenheim, A. Willsky, I. Young, *Signals and Systems*, Prentice Hall International Editions, 1983.
- [19] B. Bozkurt, L. Couvreur, T. Dutoit, *Chirp group delay analysis of speech signals*, Speech Comm., vol. 49, issue 3, pp. 159-176, 2007.
- [20] G. Sitton, C. Burrus, J. Fox, S. Treitel, *Factoring Very-High Degree Polynomials*, IEEE Signal Processing Magazine, pp. 27-42, 2003.
- [21] N. Sturmel, C. D'Alessandro, B. Doval, *A comparative evaluation of the Zeros of the Z Transform for voice source estimation*, Proc. Interspeech, 2007.
- [22] C. D'Alessandro, B. Bozkurt, B. Doval, T. Dutoit, N. Henrich, V. Tuan, N. Sturmel, *Phase-based Methods for Voice Source Analysis*, Advances in Nonlinear Speech Processing, LNCS 4885, pp. 1-27, 2008.
- [23] J. Wangrae, K. Jongkuk, B. Myung Jin, *A Study on Pitch Detection in Time-Frequency Hybrid Domain*, Lecture Notes in Computer Science, Springer Berlin, pp. 337-340, 2005.
- [24] T. Quatieri, *Minimum- and Mixed-Phase Speech Analysis/Synthesis by Adaptive Homomorphic Deconvolution*, IEEE Trans. Acoustics, Speech, and Signal Processing, vol. ASSP-27, no. 4, pp. 328-335, 1979.

- [25] K. Steiglitz, B. Dickinson, *Computation of the complex cepstrum by factorization of the z-transform*, Proc. ICASSP, vol. 2, pp. 723-726, 1977.
- [26] K. Steiglitz, B. Dickinson, *Phase unwrapping by factorization*, IEEE Trans. on ASSP, vol. 30, no. 6, pp. 984-991, 1982.
- [27] F. Nordin, T. Eriksson, *A speech spectrum distortion measure with inter-frame memory*, IEEE International Conf. on Acoustics Speech and Signal Processing, vol. 2, pp. 717-720, 2001.
- [28] K. Paliwal, B. Atal, *Efficient vector quantization of LPC parameters at 24 bits/frame*, IEEE Trans. Speech Audio Processing, vol. 1, pp. 3-14, 1993.
- [29] P. Naylor, A. Kounoudes, J. Gudnason, M. Brookes, *Estimation of glottal closure instants in voiced speech using the DYPSA algorithm*, IEEE Trans. Audio Speech Lang. Processing, vol. 15, no. 1, pp. 34-43, 2007.
- [30] T. Drugman, T. Dutoit, *Glottal Closure and Opening Instant Detection from Speech Signals*, Proc. Interspeech, 2009.
- [31] B. Doval, C. D'Alessandro, *The Spectrum of Glottal Flow Models*, Acta acustica united with acustica, vol. 92, no. 6, pp. 1026-1046, 2006.
- [32] C. Pedersen, O. Andersen, P. Dalsgaard, *ZZT-domain Immiscibility of the Opening and Closing Phases of the LF GFM under Frame Length Variations*, Proc. Interspeech, 2009.
- [33] [Online], *The Snack Sound Toolkit*, <http://www.speech.kth.se/snack/>.
- [34] M. Schroeder, M. Grice, *Expressing vocal effort in concatenative synthesis*, Proc. 15th International Conference of Phonetic Sciences, pp. 2589-2592, 2003.
- [35] P. Alku, T. Bäckström, E. Vilkman, *Normalized amplitude quotient for parametrization of the glottal flow*, the Journal of the Acoustical Society of America, vol. 112, pp. 701-710, 2002.
- [36] G. Fant, *The LF-model revisited. Transformations and frequency domain analysis*, STL-QPSR, vol. 36, no. 2-3, pp. 119-156, 1995.

- [37] D. Klatt, L. Klatt, *Analysis, synthesis and perception of voice quality variations among female and male talkers*, the Journal of the Acoustical Society of America, vol. 87, pp. 820-857, 1990.
- [38] I. Titze, J. Sundberg, *Vocal intensity in speakers and singers*, the Journal of the Acoustical Society of America, vol. 91, no. 5, pp. 2936-2946, 1992.
- [39] H. Hanson, *Individual variations in glottal characteristics of female speakers*, Proc. ICASSP, pp. 772-775, 1995.
- [40] D. Childers, C. Lee, *Vocal quality factors : analysis, synthesis, and perception*, the Journal of the Acoustical Society of America, vol. 90, no. 5, pp. 2394-2410, 1991.
- [41] M. Plumpe, T. Quatieri, D. Reynolds, *Modeling of the glottal flow derivative waveform with application to speaker identification* IEEE Trans. on Speech and Audio Processing, vol. 7, pp. 569-586, 1999.
- [42] B. Bozkurt, B. Doval, C. D'Alessandro, T. Dutoit, *A Method For Glottal Formant Frequency Estimation*, Proc. Interspeech, 2004.
- [43] T. Ananthapadmanabha, G. Fant, *Calculation of true glottal flow and its components*, Speech Commun., pp. 167-184, 1982.
- [44] T. Drugman, B. Bozkurt, T. Dutoit, *Chirp Decomposition of Speech Signals for Glottal Source Estimation*, ISCA Workshop on Non-Linear Speech Processing, 2009.