

Automatic Segmentation of Punjabi Speech Signal using Group Delay

Anupriya Sharma¹

¹ RIMT, Gobindgrah.

Received: 16 December 2012 Accepted: 2 January 2013 Published: 15 January 2013

Abstract

This paper describes the concept of automatic segmentation of continuous speech signal. The language used for segmentation is the most widely spoken language i.e. Punjabi. Like all other Indian languages, Punjabi is a syllabic language, thus syllables are selected as the basic unit of segmentation. The traditional way of representing the speech signal is in terms of features derived from short - time Fourier analysis. It is difficult to compute the phase and processing the phase function from the FT phase. By processing the derivative of the FT phase, the information in the short - time FT phase function can be extracted. This paper describes the process of automatic segmentation of speech using group delay technique. This includes segmentation of continuous Punjabi speech into syllable like units by using the high resolution properties of group delay. This group delay function is found to be a better representative of the STE function for syllable boundary detection.

Index terms— digital signal processing, speech signal, automatic speech segmentation, punjabi speech segmentation, asr, ste, syllables, units of speech.

1 Introduction

Automatic speech recognizers (ASR) are used to facilitate communication between humans and machines. So it's a machine which understands human and the words spoken by them. The process of segmentation is one of the most important phases in the automatic recognition of speech. There are various units of speech into which it can be segmented, but syllables are found to be one of the most efficient units for automatic speech segmentation. The characteristics features of speech can be expressed by using STE and ZCR. The STE function also known as Short Term Energy function is known to be the better representative of speech segment boundaries. By computing the shorttime Fourier analysis information in the speech signal can be extracted. But due to difficulty in computing the phase and also in processing the phase function over the past few decades the features of the FT phase were not exploited fully. By processing the derivative of the FT phase, the information in the short-time FT phase function can be extracted. There are various units of speech. The syllables are found to be the most suitable unit for automatic speech segmentation. A single component in the syllable is called as nucleus. The nucleus is found to be vowel while the onset and coda are usually consonantal in form. The energy peak in the nucleus region can be viewed as the syllable; the consonants can be viewed as the valleys at both the ends. Many languages been spoken around the world possess a syllabic structure [10]. Mostly the syllable contains two phonetic segments of type CV such as in Japanese language. In contrast, English and German possess a more highly heterogeneous syllable structure [2].

2 II.

3 Research Background a) Language Units of Speech in Punjabi

Punjabi is an Aryan language that is spoken by more than hundred million people those are inhabitants of the historical Punjab region (in north western India and Pakistan) and in the Diaspora, particularly Britain, Canada, North America, East Africa and Australasia [8].

Like other Indian languages the Punjabi language also contains segmental phonemes. The three basic units into which the speech can be segmented are: Words, Phonemes and Syllables. The syllable is the most important and widely used unit for automatic speech segmentation. Punjabi is a syllabic language thus syllables are selected as the basic units for segmentation.

4 b) Syllables as Basic unit of speech

Aksharas is the basic units of the writing system. An Akshara is an orthographic representation of a speech sound in an Indian language. Basically they are syllabic in nature; the typical forms of akshara are V, CV, CCV and CCCV type, where C and V are consonant vowel respectively [9]. There are thirty eight consonants in Punjabi language. Where ten are non-nasal and ten are nasal vowels. Vowels can appear alone but consonants can only appear with vowels. The number of nasal vowels is same as non-nasal vowels and is represented by Bindi or Tippi over the Non-Nasal Vowels. Following is the list of consonants in Punjabi language:

5 Three State Representation of Speech

The continuous speech signals composed of two elements one includes the speech information, and the other carries noise or silent sections. The verbal part of the speech can be further divided into two categories: voiced and unvoiced speech. Moment the air from the lungs passes through the larynx voiced sound is produced. With the passage of air directly through the vocal tract formations the unvoiced speech sounds are produced. The speech production process is incomplete without the detection of voiced and unvoiced speech that is separated by a silence region. In case of silence region no excitation is supplied to the vocal tract and thus, no speech is produced. A regular speech is incomplete inaccurate without silence region. It helps to make the speech understandable [3].

IV.

6 Characterization OF Speech

In order to segment continuous speech it is required to check its basic content, whether the signal is voiced or unvoiced. The two characteristics features of voice are the zero crossing rate (ZCR) and short term energy (STE) [13].

7 a) Zero Crossing Rate

The rate at which the signal crosses zero provides the information regarding its (source of creation) i.e. zero crossing rate. Unvoiced speech has higher zero crossing rate. Whereas in case of voiced speech the zero crossing rate is low. Thus, the amplitude of unvoiced segments is lower than that of the voiced segments. ZCR can be defined as: The STE can be defined as follows:

(2)

The STE of voiced signal is always much greater than that of unvoiced signals. In a speech signal where there are voiced signal its STE will be high, the peaks in the signal represents nucleus that is denoted as vowel where as the valleys at both the ends represents the coda.

8 SEGMENTATION OF SPEECH

The syllable is composed of three parts, the onset, rime (nucleus) and coda. The rime also known as nuclei, where as the onset and coda consist of consonants. The high energy regions are represented by the nuclei where as the valleys at both ends corresponds to syllable boundaries. The vowel region corresponds to much higher energy region compared to that of a consonant region [9]. In case of spontaneous speech, the definition of a syllable in terms of short-term energy function is suitable for almost all the languages.

Due to local energy fluctuations the STE function alone cannot be directly used to perform segmentation. Techniques such as fixed or even adaptive threshold will not work when the energy variation across the signal is quite high [1].

To overcome the problems of local energy fluctuations, the STE function should be smoothed. The information in speech signals can be represented in terms of features derived from short-time Fourier analysis. The information in the short-time FT phase function can be extracted by computing the group delay function [9]. $H(?) = H1(?) ? H2(?), (3)$

group delay function can be represented as $h(?) = ??(\arg(H(?))) \rightarrow ?? = ?h1(?) + ?h2(?). (4)$

The equation (1) shows the multiplicative property of magnitude spectra where as equation (??) is in group delay domain it shows an addition. The group delay spectrum has been found better due to its additive. It

was observed that in case of the magnitude spectra the peaks are clearly visible, but when the two poles are combined together the peaks are not resolved. The research shows the disadvantage of multiplicative property of magnitude spectra. In case of group delay spectra the peaks and valleys are better resolved when the signal is in minimum phase [2].

For any syllable, the STE function of the voiced region, the energy is quite high and diminishes at the ends, representing the consonants, due to which local energy fluctuations. If these local variations are smoothed, then the minima at both ends of a voiced region correspond to syllable boundaries [9].

9 The algorithm for group delay based segmentation

Step 1 -Let $x[n]$ be continuous speech signal.

Step 2 -Compute N , the length(x) of the input signal.

Step 3 -Calculate the STE function $E[m]$, where $m=1,2,\dots,M$ is the number of frames.

Step 4 -Inverse the STE i.e $E(i) = 1/E(m)$

Step 5 -Compute the IFFT of $E(i)$, It gives the magnitude of the input signal in form of complex function i.e. $a+ib$.

Step 6 -The phase angle is computed from the above values, i.e. $\theta = \tan^{-1}(b/a)$.

Step 7 -Compute the negative derivative of Fourier transformation i.e. the group delay function.

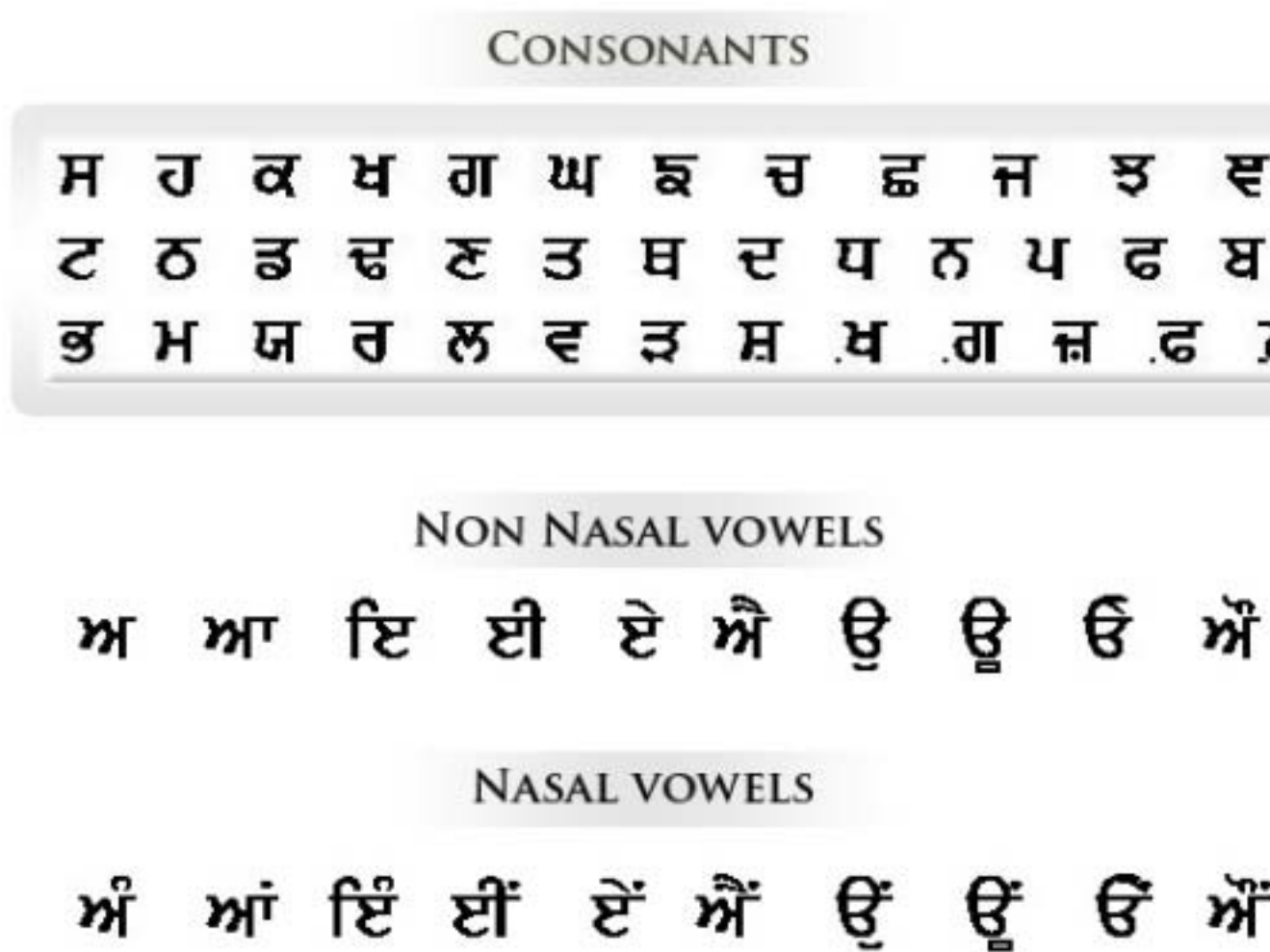
10 Results and discussions

The technique of automatic segmentation is applied on the continuous Punjabi speech. The method was implemented in Matlab. The group delay algorithm is applied to segment the continuous Punjabi speech waveform. The following sentence is given as an input to the system.



20131

Figure 1: T © 2013 Fi gure 1 :



2

Figure 2: Figure 2 :

3

Syllable	Syllable Description	Word	Segments	Word	Segments
V	Vowel	ਉ	ਉ	ਈ	ਈ
VC	Vowel-consonant	ਉਡ	ਉ+ ਡ	ਇਸ	ਇ+ ਸ
CV	Consonant-Vowel	ਜਾ	ਜ + ਆ	ਗਾ	ਗ + ਆ
VCC	Vowel-consonant-consonant	ਉਤਰ	ਉ+ ਤ+ਰ	ਅੰਦਰ	ਅੰ+ਦ+ਰ
CVC	Consonant-vowel-consonant	ਰਾਤ	ਰ + ਆ + ਤ	ਬਾਤ	ਬ + ਆ + ਤ
CVCC	Consonant-vowel-consonant-consonant	ਜੋਤਸ਼	ਜ + ਔ + ਤ + ਸ਼	ਪੂਰਬ	ਪ + ਊ + ਰ + ਬ
CCVC	Consonant-consonant-vowel-consonant	ਤਰੇਲ	ਤ + ਰ + ਏ + ਲ	ਸਵੇਰ	ਸ + ਵ + ਏ + ਰ

Figure 3: Figure 3 :

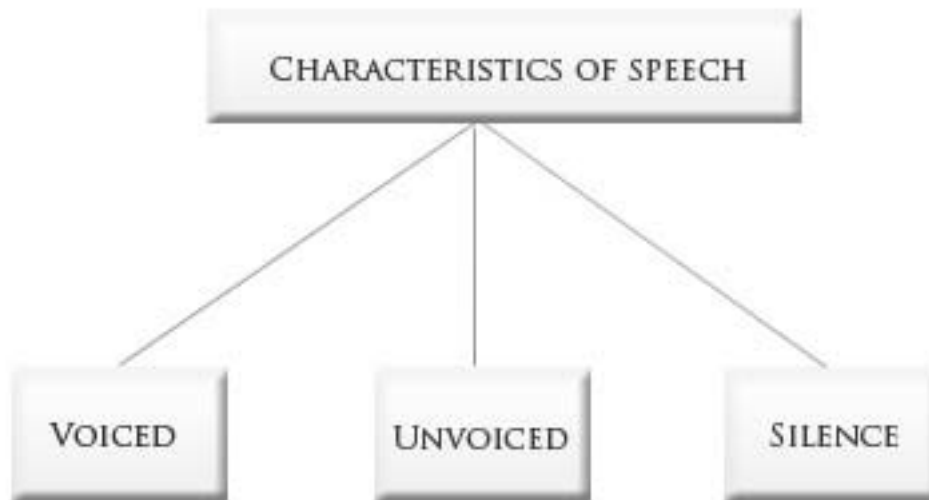


Figure 4:

$$Z_n = \sum_{m=-\infty}^{\infty} |\text{sgn}[x(m)] - \text{sgn}[x(m-1)]| w(n-m)$$

Figure 5: C

$$\text{sgn}[x(n)] = \begin{cases} 1, x(n) \geq 0 \\ -1, x(n) < 0 \end{cases}$$

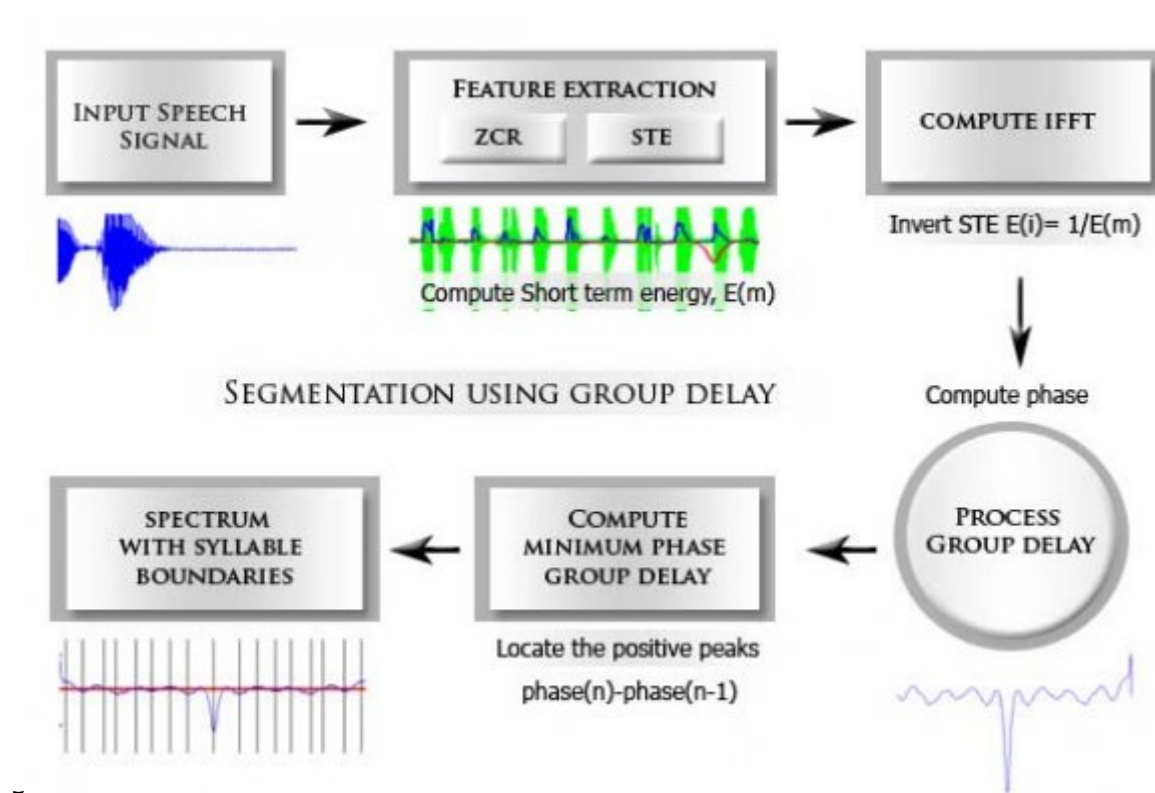
8

Figure 6: Step 8 -

$$E_n = \sum_{m=-\infty}^{\infty} [x(m)w(n-m)]^2$$

4

Figure 7: Figure 4 :



5

Figure 8: Figure 5 :

62 ਅੰਮ੍ਰਿਤਸਰ ਸਿੱਖਾਂ ਦਾ ਸਭ ਤੋਂ ਉੱਚਾ ਧਾਰਮਿਕ ਸਥਾਨ ਹੈ

Figure 9: Figure 6 : 2 C

Sentence	Onset	Offset	Duration
AMimR	0.576	1.664	1.088
qsr	1.664	2.88	1.216
isW	2.88	3.904	1.024
KF	3.904	5.056	1.152

Figure 10:

-
- [Gauvian (1986)] ‘A syllable-based isolated word recognition experiment’. J L Gauvian . *IEEE International Conf. on ICASSP’86*, April 1986. 11 p. .
- [Nishi and Parminder (2010)] ‘Automatic Segmentation of Wave File’. S Nishi , S Parminder . *International Journal of Computer Science & Communication* July-December 2010. 1 (2) p. .
- [G Lakshmi Sar Ada (2009)] ‘Automatic transcription of continuous speech into syllable-like units for Indian languages’. G Lakshmi Sar Ada . *Sadhana*, April 2009. 34 p. .
- [Hema and Ashwin (2009)] *Building Unit Selection Speech Synthesis in Indian Languages*, A Hema , B Ashwin . 2009. (An Initiative by an Indian Consortium)
- [Hu et al. (2008)] *Center for Spoken Language Understanding, Oregon Graduate Institute of Science and Technology*, Zhihong Hu , Johan Schalkwyk , Etienne Barnard , Ronald Cole . September 2008. p. . (Speech Recognition Using Syllable-Like Units)
- [Parminder and Gurpreet (2010)] ‘Corpus Based Statistical Analysis of Punjabi Syllables for Preparation of Punjabi Speech Database’. S Parminder , L Gurpreet . *International Journal of Intelligent Computing Research* June 2010. 1 (3) .
- [Hema and Yegnanarayan (2011)] ‘Group delay functions and its applications in speech technology’. A Hema , B Yegnanarayan . *Sadhana*, October 2011. 36 p. .
- [Amanpreet and Tarandeep (2010)] ‘Segmentation of Continuous Punjabi Speech Signal into Syllables’. K Amanpreet , S Tarandeep . *the Proceedings of the World Congress on Engineering and Computer Science*, (San Francisco, USA) 2010. October 20-22, 2010. I. (WCECS 2010)
- [Nagarajan (2003)] ‘Segmentation of speech into syllable-like units’. T Nagarajan . *Eurospeech Sixth biennial conference of signal processing*, (Geneva) 2003.
- [Bachu et al. (2010)] ‘Separation of Voiced and Unvoiced using Zero crossing rate and Energy of the Speech Signal’. R G Bachu , S Kopparthi , B Adapa , B D Barkana . *Electrical –Engineering Department School of Engineering* March 2010. 2012. 7340 p. . University of Bridgeport
- [Mikael and Marcus (2002)] *Speech Recognition using Hidden Markov Model, Performance evaluation in noisy environment*, E Mikael , Marcus . March 2002. Department of telecommunications and engineering, Blekinge Institute of Technology (Degree of master of science in Electrical Engineering)
- [Nagarajan and Murthy (2004)] ‘Subband-Based Group Delay Segmentation of Spontaneous Speech into Syllable-Like Units’. T Nagarajan , H A Murthy . *Eurasip Journal on Applied Signal Processing* 2004. Hindawi Publishing Corporation. 17 p. .
- [Fujimura (1975)] ‘Syllable as a unit of speech recognition’. O Fujimura . *IEEE Trans. Acoust., Speech, Signal Processing* February 1975. 23 p. .
- [Pradeep (2004)] ‘Text-to-Speech Synthesis for Punjabi Language’. G Pradeep . *the proceedings of science direct*, 2004. 42 p. . Thesis degree of Master of Engineering in Software Engineering submitted in Computer Science and Engineering Department of Thapar using minimum phase group delay functions
- [Shneiderman (2000)] ‘The Limits of Speech Recognition’. B Shneiderman . *Communications of the ACM* September 2000. 43 (9) p. .