

IMPROVED METHODS OF HEART RATE MONITORING USING ADAPTIVE FILTERS

Alampreet Chahal¹, Gautam Kaushal², Dr. Rajbir Kaur³, Silki Baghla⁴

¹Department of Electronics & Communication, Punjabi University, Patiala

²Department of Electronics & Communication, Punjabi University, Patiala

³Department of Electronics & Communication, Punjabi University, Patiala

⁴Department of Electronics & Communication, JCDD college of engineering, Sirsa

Abstract- This paper presents a relation between heart rate and vowel speech signal. The proposed method is based on modeling the relationship between speech production of vowel speech signals and heart activities of human. The Adaptive filter is used to manipulate signals to reject unwanted characteristics for accurately monitoring the heart rate. Simulation results show that the proposed method provides more accuracy as compared with existing methods. There is linear relationship between results obtained from proposed method and clinical method.

Keywords: Heart Rate, Vowel speech signal, Adaptive filter, Electrocardiogram.

I. INTRODUCTION

Heart rate is measured by detecting arterial pulsation. Heart rate indicates that the total number of times our heart contracts and relaxes per minute and is expressed as the number of beats per minute (bpm). The heart rate of a healthy adult at rest is around 72 bpm. Babies have a much higher heart rate at around 120 bpm, while older children have heart rates at around 90 bpm. The heart rate rises gradually during exercises and returns slowly to the rest value after exercise [1]. A device is used to monitor and record the heart rate in real time is referred as heart rate monitoring. The earlier models of heart rate monitors use electrode leads that are attached to the chest. Modern heart rate monitors use chest strap transmitter, a mobile phone and a wrist receiver. The detection of heart rate from human voices was done based on the modeling of vowel speech signals. Heart Rate also depends on the body's need to absorb oxygen and excrete carbon-dioxide. The heart is a muscular pump made up of four chambers. The two upper chambers are called atria, and the two lower chambers are called ventricles [2].

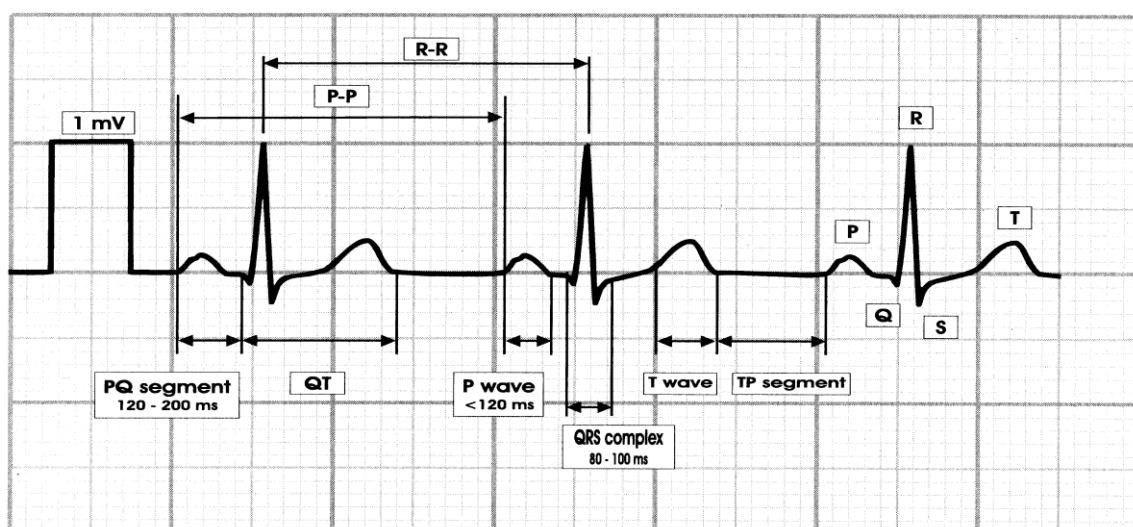


Figure1 Typical ECG Waveform [13]

The electrical activity of the heart is generally sensed by monitoring electrodes placed on the skin surface. The electrical signal is very small (normally 0.0001 to 0.003 volt). These signals are within the frequency range of 0.05 to 100 Hertz (Hz.) or cycles per second. A typical ECG tracing of a normal heartbeat (or cardiac cycle) consists of a P wave, a QRS complex and a T wave. A small U wave is normally visible in 50 to 75% of ECGs. The baseline voltage of the electrocardiogram is known as the isoelectric line. Typically the isoelectric line is measured as the portion of the tracing following the T wave and preceding the next P wave. A typical scalar electrocardiographic lead is shown in Figure 1, where the significant features of the waveform are the P, Q, R, S, and T waves, the duration of each wave, and certain time intervals such as the P-R, S-T, and Q-T intervals. ECG signal is periodic with fundamental frequency determined by the heartbeat. ECG signal is a mixture of triangular and sinusoidal wave forms. Each significant feature of ECG signal can be represented by shifted and scaled versions one of these waveforms as shown below [3].

- QRS, Q and S portions of ECG signal can be represented by triangular waveforms.
- P, T and U portions can be represented by triangular waveforms.

II. LITERATURE REVIEW

In previous research, several algorithms have been proposed to estimate heart rate from vowel speech signal. In [4], [5], and [6] present Heart Rate methods by using machine learning algorithms, such as the support vector machine (SVM). However, the explicit relationship between Heart Rate and voice has not revealed, and it is unclear which features of voice were conducive to estimate Heart Rate. For example, the Support Vector Machine does not require explicit features of voice to estimate HR because learning and regression are conducted in unknown feature space. In [7], the fundamental concepts of Linear Predictive Coding (LPC) were formulated for estimation of the vocal tract response from speech waveforms. An algorithm is used to find the best match between the input pattern and the reference pattern is derived. A sequential decision procedure is used to reduce the amount of computation in dynamic programming algorithm. The recognition time is about 22 times real time. In [8], Heart Rate is roughly estimated by image processing for 2D images of speech spectrogram. In [9] the voice signal is considered to contain biometric data reflective of the physical condition of the heart e.g. heart rate index. There is an implicit relationship between heart rates and human voice. In [10], authors proposed a method for the detection of heart rate from human speech based on the modeling of vowel speech signals. The non contact method for the detection of heart rate from human speech is based on the modeling the relationship between speech production of vowel speech signals and heart activities for humans. It uses Short Time Fourier Transform to estimate heart rate from vowel speech signals. Hidden Markov Model (HMM) [11] is a natural and highly robust statistical methodology for automatic speech recognition. The model parameters of the HMM are essence in describing the behavior of the utterance of the speech segments. Many successful heuristic algorithms are developed to optimize the model parameters in order to best describe the trained observation sequences. The objective of this paper is to develop an efficient speech recognition algorithm with the existing system following HMM algorithm. In [12] the several methods are used Mel Frequency Cepstral Coefficient (MFCC), Linear Predictive Coding (LPC), Hidden Markov model (HMM) and Artificial Neural Networks (ANN) to identify a straight forward and effective method for voice signal extraction. In [13] there is a strong correlation between the human speech and the heart rate. The proposed method provides more accuracy as compared to existing method.

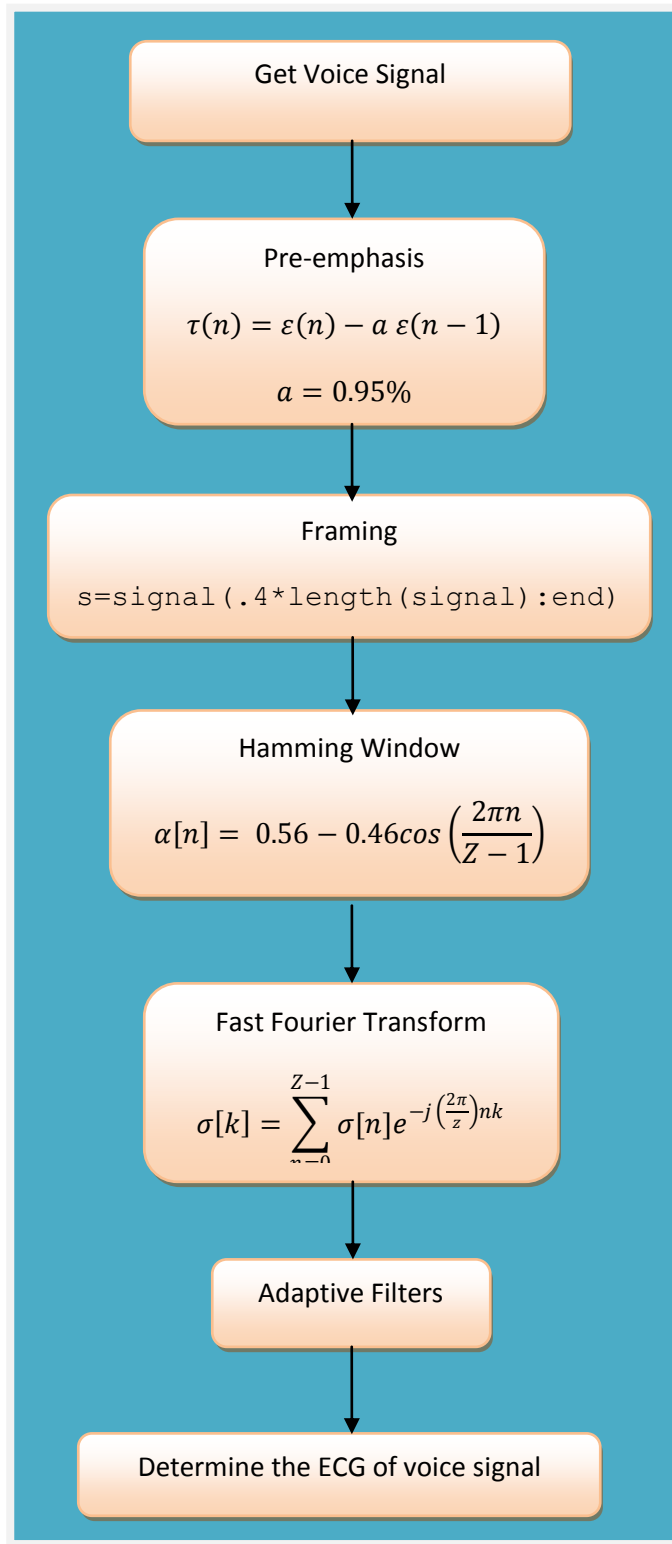


Figure2 Phases of speech signal

III. PROBLEM FORMULATION

Contactless heart rate monitoring methods are required for monitoring conditions of heart patients. Several techniques used for voice recognition such as Mel Frequency Cepstral Coefficient(MFCC), Linear Predictive Coding(LPC), Artificial Neural Networks(ANN) and Hidden Markov Model(HMM) has been used in literature to correlate the features of voice with that of human heart rate. The methods are based on human emotions which cannot give accuracy comparable to traditional ECG. Moreover these techniques do not consider artificial heart or transplanted hearts. Thus techniques are required which can consider human voice as indicator of heart rate and are applicable for artificial as well as transplanted hearts.

IV. METHODOLOGY OF PROPOSED WORK

The various steps of voice signal feature extractions are as follow:

3.1 Data Collection

The vowel speech signal and the corresponding ECG are recorded concurrently. Totally, 20 subjects participated in the data collection process. During data collection, the microphone was kept closer to the mouth to reduce noise and increase the sound level of the speech signal. The device has 30 seconds recording capacity in its easy mode of capacity. All the analysis performed in this paper and takes into account the dynamic nature of heart rate so as to reflect in real time applications [13].

3.2 Feature Extraction

The extraction of features from the vowel speech signal is important task to understand the human speech signal. Feature distances were extracted by applying Pre-emphasis, Framing, Hamming Window, Fast Fourier Transform and adaptive filtering. The identification process includes five important steps: pre-emphasis, framing, hamming window, Fast Fourier transform and adaptive filter [13].

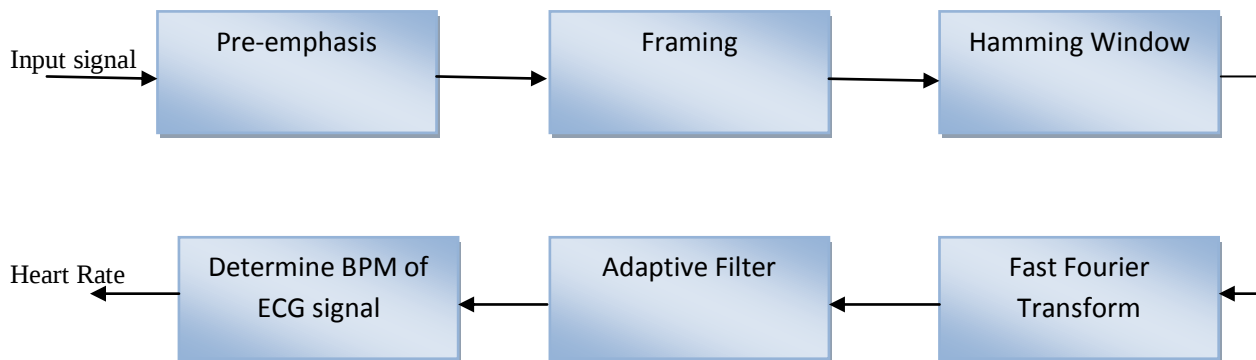


Figure3 Feature Extraction Steps [13]

3.2.1 Pre-emphasis:

This step processes the passing of signal through a filter which emphasizes higher frequencies. This will increase the energy of signal at higher frequency because the low frequency band is occupied by sounds which are useless for speech recognition. The signals with higher modulation frequencies have lower Signal to Noise Ratio [13].

$$\tau(n) = \varepsilon(n) - a \varepsilon(n-1), \quad (1)$$

$\tau(n)$ = output signal, $\varepsilon(n)$ = input signal and n = "time" index of discrete time signal. Let consider $a = 0.95$, which make 95% of any one sample is presumed to originate from previous sample.

3.2.2 Framing:

Framing is the process of segmentation of speech samples. The speech signal is divided into frames of 20ms which corresponds to Z samples. Adjacent frames are being separated by N ($N < M$) values are used. $N=100$ and $M=256$. Where M = number of samples per frames and N = adjacent frames [13].

3.2.3 Hamming Window:

Hamming window is used to reduce the effect of leakage for better representation of the frequency spectrum of the speech signal. It helps in reducing the spectral artifacts of the speech signal. The window function is denoted as $\alpha(n)$ [13].

$$\alpha[n] = 0.56 - 0.46\cos\left(\frac{2\pi n}{Z-1}\right), \quad 0 \leq n \leq Z-1 \quad (2)$$

$$\tau(n) = \varepsilon(n) * \alpha(n) \quad (3)$$

Where Z = number of samples in each frame, $\alpha(n)$ = hamming window, $\tau(n)$ = output signal, $\varepsilon(n)$ = input signal.

3.2.4 Fast Fourier Transform:

To convert each frame of Z samples from time domain into frequency domain, FFT is used. In this work, The Fourier transform is used to convert the convolution of the glottal pulse $G(n)$ and vocal tract impulse response $I(n)$ in the time domain [13].

$$\sigma[k] = \sum_{n=0}^{Z-1} \sigma[n] e^{-j\left(\frac{2\pi}{Z}\right)nk}, \quad k = \{0, 1, 2, \dots, Z-1\} \quad (4)$$

Where $\sigma[k]$ = Fourier transform of input signal $\sigma[n]$, Z = number of samples per frame

3.2.5 Adaptive filter:

Adaptive filter is used to manipulate the signals to reject unwanted characteristics. Noise, Interference, erroneous frequencies are undesirable characteristics that affect data in signals. The least-mean squares (LMS) algorithm is simple to implement and aids in identifying areas of improvement for the characteristics of the adaptive filter. After areas of improvement have been identified, more sophisticated algorithms can be substituted in place of the LMS algorithm.

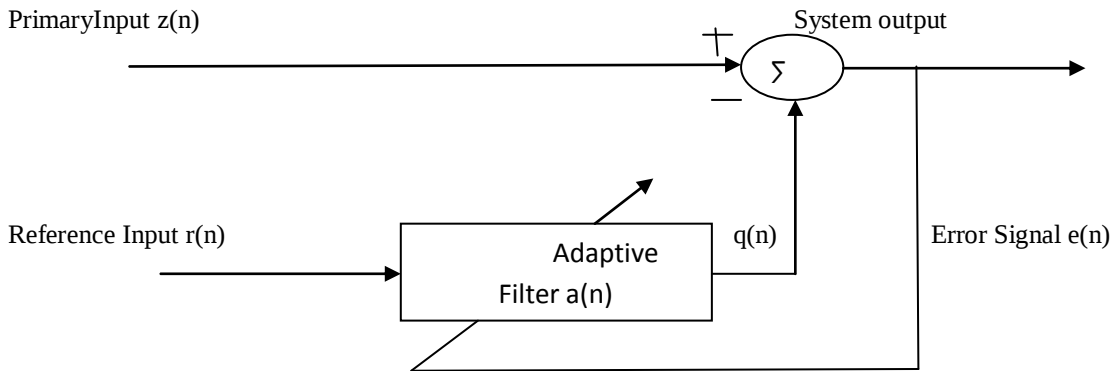


Figure4 Adaptive Filter Block Diagram [13]

Figure.3 represents block diagram of adaptive filters. The primary input $z(n)$ is the incoming signal that contains the input signal $\varepsilon(n)$ and pulse noise $q(n)$. However, the noise is uncorrelated with the original signal. The reference input $r(n)$ is measured separately and contains noise that is correlated with the noise $r(n)$ from the primary input. The noise from the reference input passes through an adaptive filter that estimates the noise of the original signal is known as $q(n)$. $q(n)$ is estimated by the adaptive filter coefficients and is driven by the error signal $e(n)$. The error signal is the difference between the desired signal and the incoming signal is the estimated signal output [13].

3.3 HEART RATE DETECTION:

A human heart rate is in the range of 60-100 bpm, but this may change according to age, sex and size of the person. The QRS complex represents the depolarization of ventricles. The heart rate is calculated as the time interval between the two QRS

complexes per unit time. The duration of QRS complex is generally 0.06-0.1s. The heart rate is obtained from recorded ECGs using 1500 rule. Count the number of small squares between two neighboring R-waves and divide that number with 1500. The obtained value will be the heart rate, H of the respective speech [13].

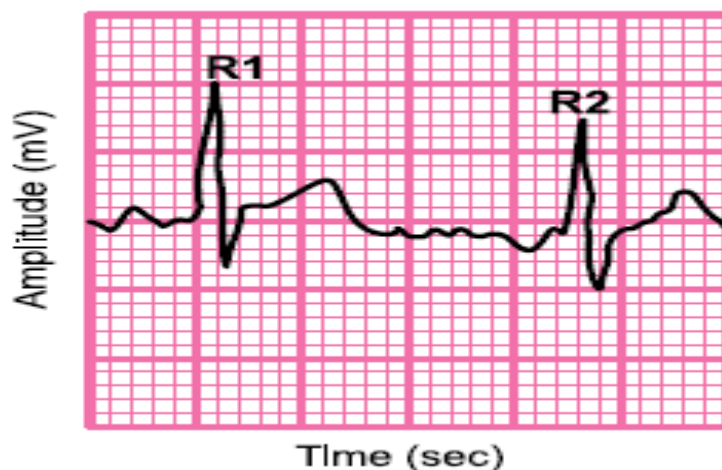


Figure5 Electrocardiograph [13]

$$\text{Heart rate} = \frac{1500}{\text{number of small boxes between two } R - R \text{ intervals}} \quad (5)$$

V. Results and discussion

Several techniques are proposed to extract speech features. Now showing the principle of feature extraction of the speech signal in MATLAB, we need to correlate it to the ECG signal. The average feature difference of the Adaptive filter and the heart rate measured from the recorded ECG are sampled on 20 different samples. The distance features drawn from the adaptive filtering and the heart rate extracted from the ECG is correlated by the samples to predict the vowel speech. The input speech signal was recorded in the ECG laboratory. Figure6 shows the variation of sound signal with respect to time.

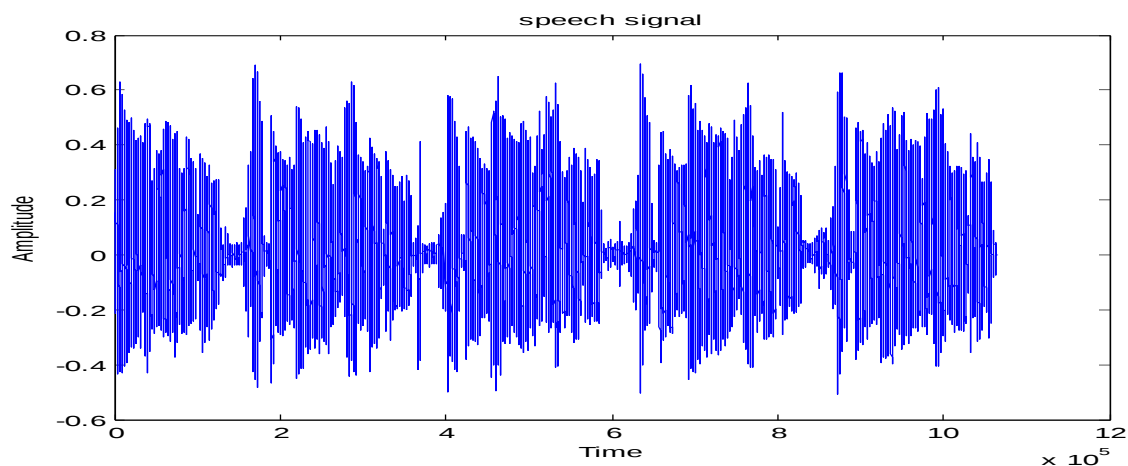


Figure6 Speech Signal

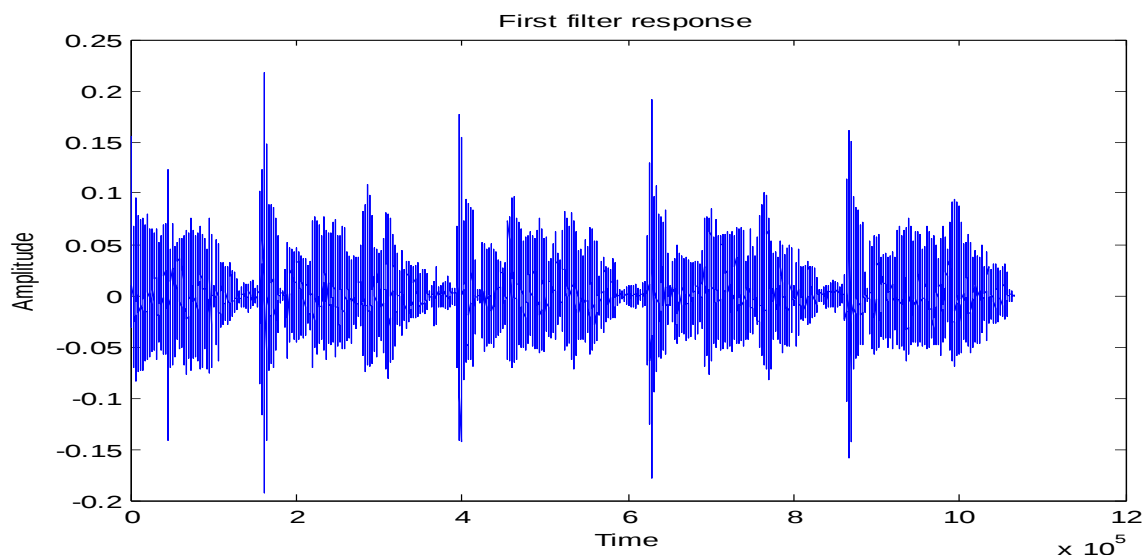


Figure7 First Filter Response

Figure7 shows the input speech signal is passing through a digital filter which emphasizes the higher frequency.

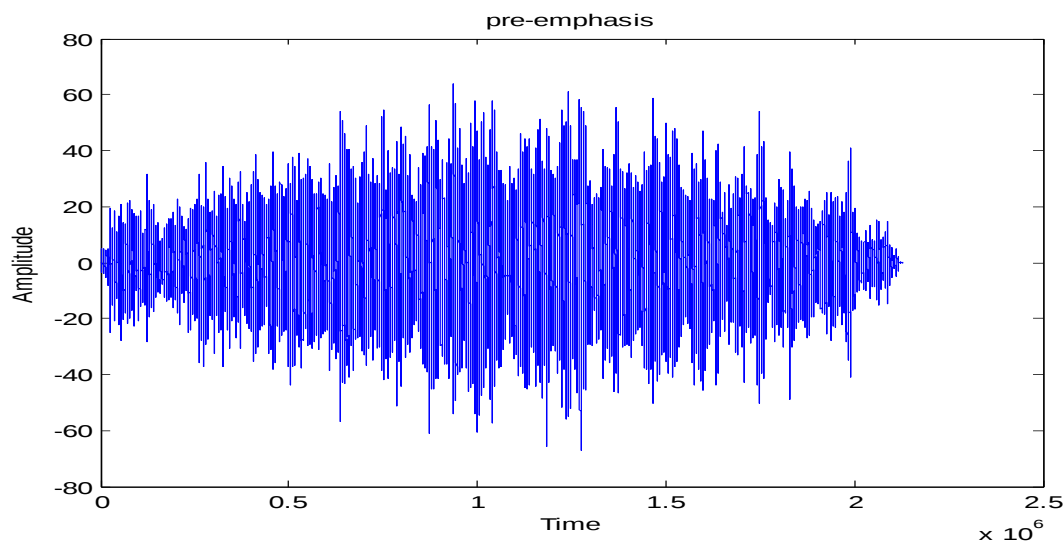
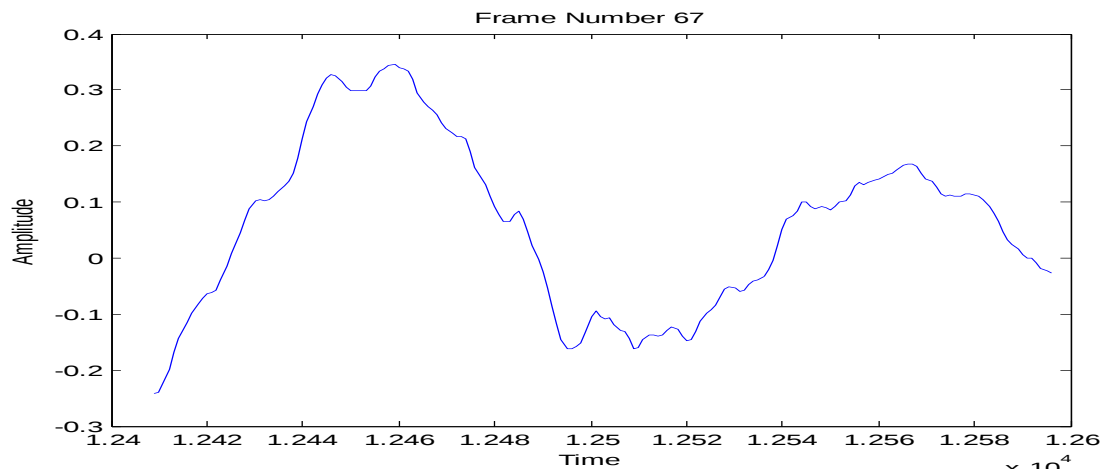


Figure8 Pre-emphasis after filtering

Figure8 shows that the Pre-emphasis after filtering will increase the energy of signal at higher frequency.



Frame Number 67

Figure9

The framing is used to divide the voice samples into frames of M samples ($M=256$). Figure 9 shows the results of frame number 67. We take the output of any of the frame number 100 to 256.

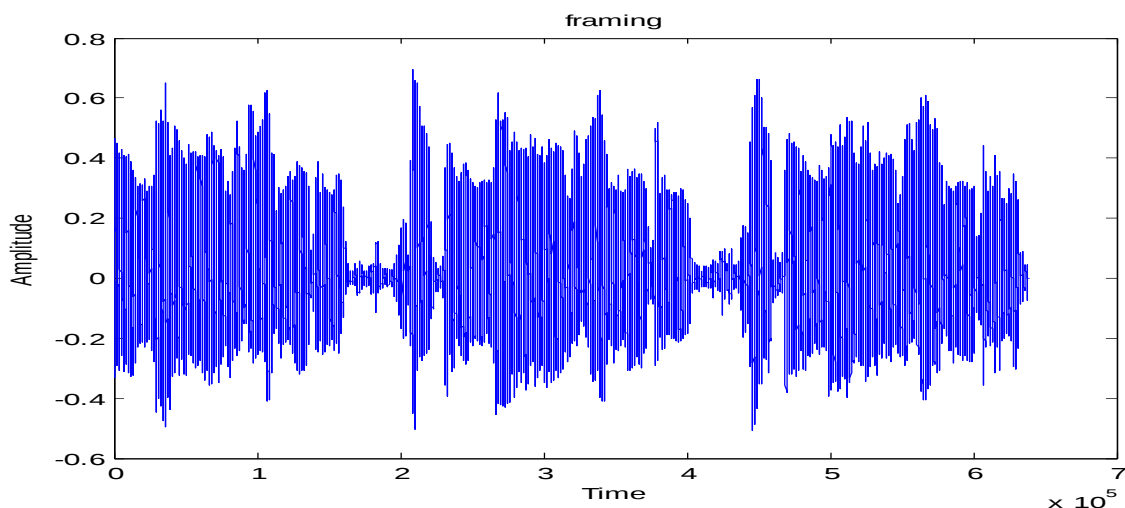


Figure10 Framing

The output shows that the framing of a frame number 150. The framing is used to remove the silences from a speech signal and create a new signal which does not contain silent signal.

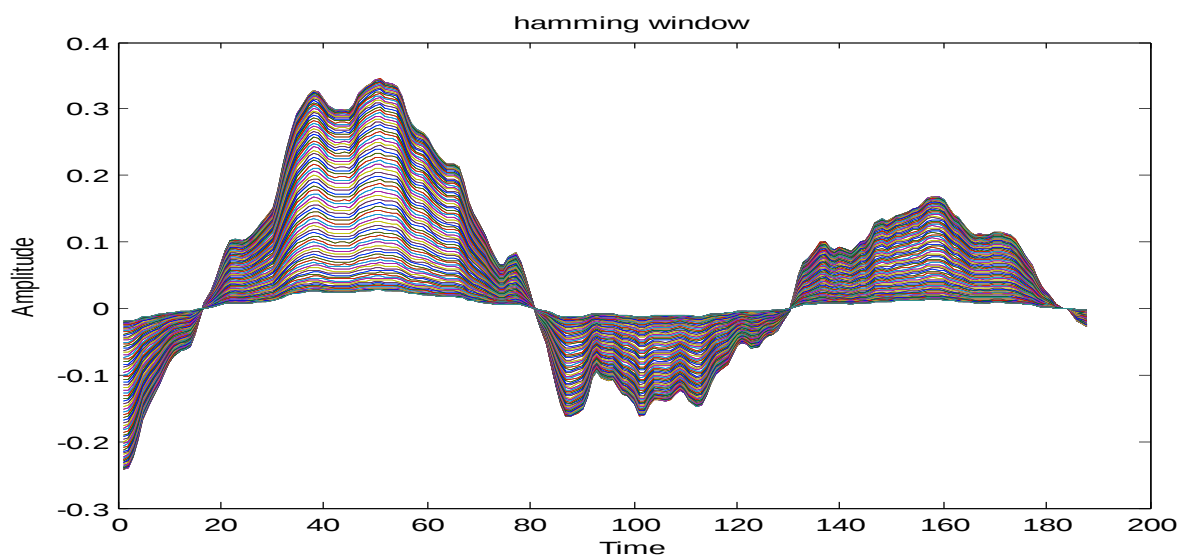


Figure11 Hamming Window

The effect of leakage is reduced for better representation of the frequency spectrum of the speech signal in figure 10. The frames obtained are multiplied with the window function $a[n]$ to reduce the discontinuities of the speech signal in the time domain.

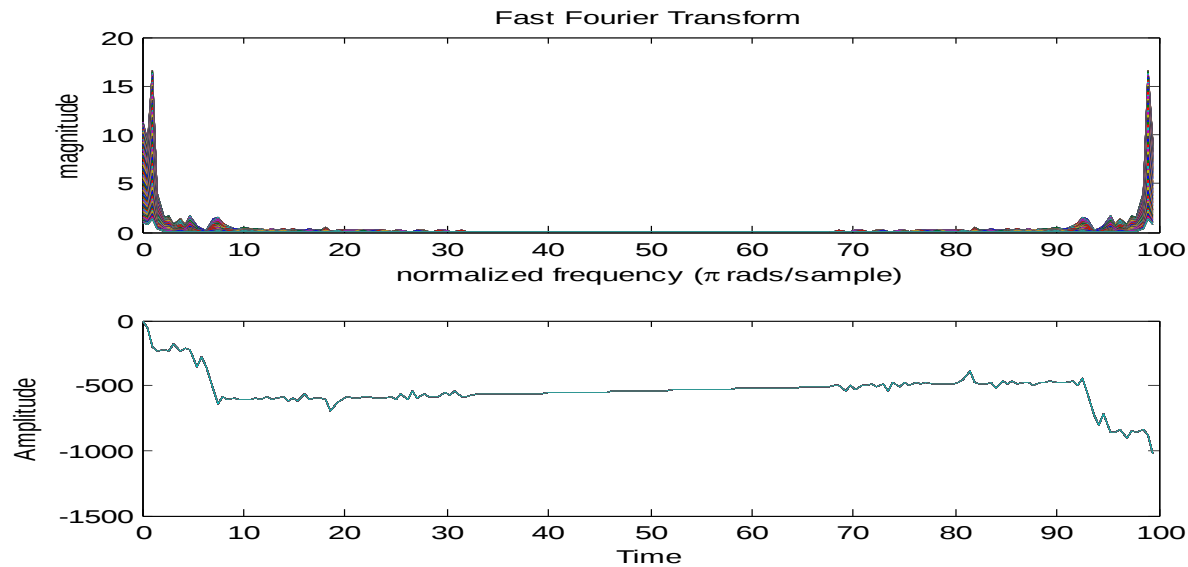


Figure12 Fast Fourier Transform of Speech Signal

The Fast Fourier Transform algorithm to compute the Discrete Fourier transform produces exactly the same result as evaluating the DFT; the most important difference is that an FFT is much faster. In the highest peak shows the sinusoidal signal and the lowest peak shows the noise. A broad peak resulting in poorly determined frequency and inaccurate amplitude. These waveforms were generated by an inexpensive function generator, which accounts for the noise present in the spectrum.

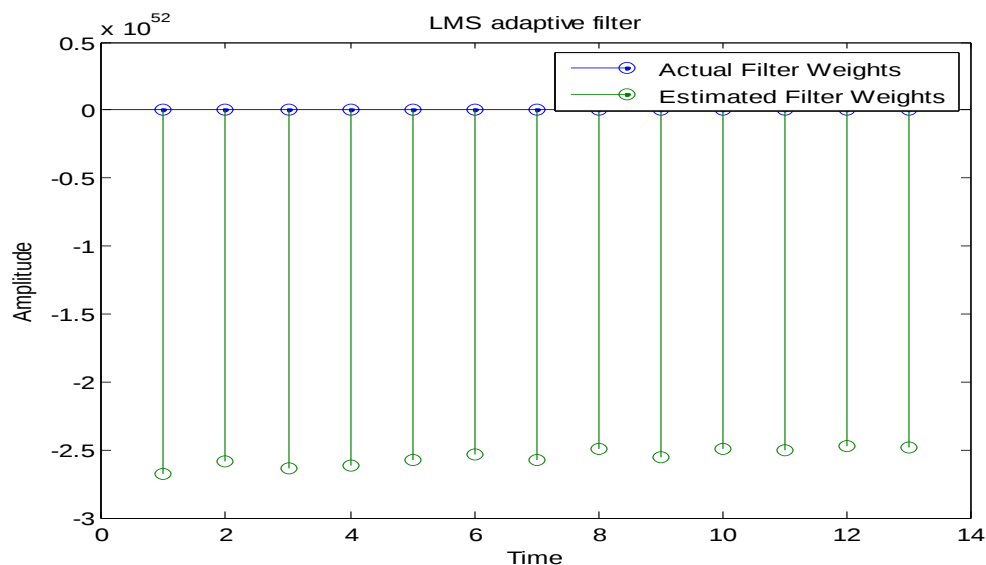


Figure13 LMS Adaptive Filter

In figure13, the filter is used to manipulate the signals to reject unwanted characteristics like Noise, Interference, erroneous frequencies that affect the signals. To improve the convergence performance of the LMS algorithm, the normalized variant uses an adaptive step size based on the signal power. As the input signal power changes, the algorithm calculates the input power and adjusts the step size to maintain an appropriate value. Thus the step size changes with time.

The ECG signal is the signal that is measure the electrical activity of the heart. The large spike is very prominent features that we can use so every heart beat contains large prosperity of spike our electro activity. The result shows that there is strong correlation between the heart rate and the human speech. Since, we have recorded the vowel speech signals of different patients and ECG signal simultaneously, the results. Moreover the voice samples are recorded in clinical rooms which are not sound proof. Thus the difference between heart rate extracted from voice samples and clinical ECG signal is greater. In table.1 we have calculated the error percentage of the heart rate from clinical ECG and the heart rate from proposed method. The results obtained from the clinical heart rate and the proposed heart rate is shown 65.84% accuracy according to this formula.

$$ErrorPercentage = \left(\frac{HR(Clinical) - HR(Proposed)}{HR(Clinical)} \right) \times 100$$

The curve plotted between the Heart Rate and the Beat Count for clinical as well as proposed method is shown in Figure 14. It can be concluded that the variation of the Heart Rate (proposed) is linear to the Heart Rate (clinical).

Table 1 Error Percentage of Different Samples

S.No.	Voice samples	Heart rate (from clinical ECG)	Heart rate (Proposed Method)	Error (Percentage)
1.	Voice1	125	130	4
2.	Voice2	98	127	29
3.	Voice3	78	130	66
4.	Voice4	107	127	18
5.	Voice5	83	130	56
6.	Voice6	107	134	25
7.	Voice7	88	130	47
8.	Voice8	88	127	44
9.	Voice9	78	130	66
10.	Voice10	80	130	62
11.	Voice11	83	130	56
12.	Voice12	125	130	4
13.	Voice13	78	125	60
14.	Voice14	107	134	25
15.	Voice15	78	130	66
16.	Voice16	75	127	18
17.	Voice17	115	130	13
18.	Voice18	105	130	23
19.	Voice19	89	134	50
20.	Voice20	93	130	39
21.	Voice 21	125	130	4
22.	Voice 22	105	130	23
23.	Voice 23	107	127	18
24.	Voice 24	107	134	25
25.	Voice 25	115	130	13
TOTAL				854

$$\text{Percentage error} = \frac{854}{25} = 34.16\%$$

$$\text{Accuracy} = (100 - 34.16)\% = 65.84\%$$

Table 2 Performance Comparison

S. No.	Number of Samples	Percentage Error	Accuracy
1.	20	41.1%	58.9%
2.	25	34.16%	65.84%

We have used 20 voice samples in my previous work. Now we have taking 25 voice samples for better accuracy. The proposed method shows an accuracy 65.84% which is better than a previous method.

VI. CONCLUSION

In this work samples are collected from patients present for ECG in clinical laboratory. The sound signals have been recorded while recording their ECG. Thus, the observations are taken in real environment and results can be correlated with readings of clinical ECG. Simulation results show a linear relationship with clinical data and hence proved effectiveness of proposed method. Though the proposed method shows a limited accuracy of 65.84% which is better than previous methods. Heart rate monitoring using adaptive filters shows a strong correlation between the human speech and the heart rate. The simulation results shows accuracy 58.9% [13]. Filtered classifier is a class for running an arbitrary classifier on data that has been passed through an arbitrary filter. Filtered classifier used Meta filtered classifier gives accuracy 33.33% [14]. J48 is a classifier which is implemented by C4.5 algorithm. J48 is an open source java implementation of the C4.5 algorithm in the Weka data mining tool. The J48 classifier gives an accuracy of 39.24% [15]. The accuracy can be increased by taking the voice samples in sound proof rooms etc.

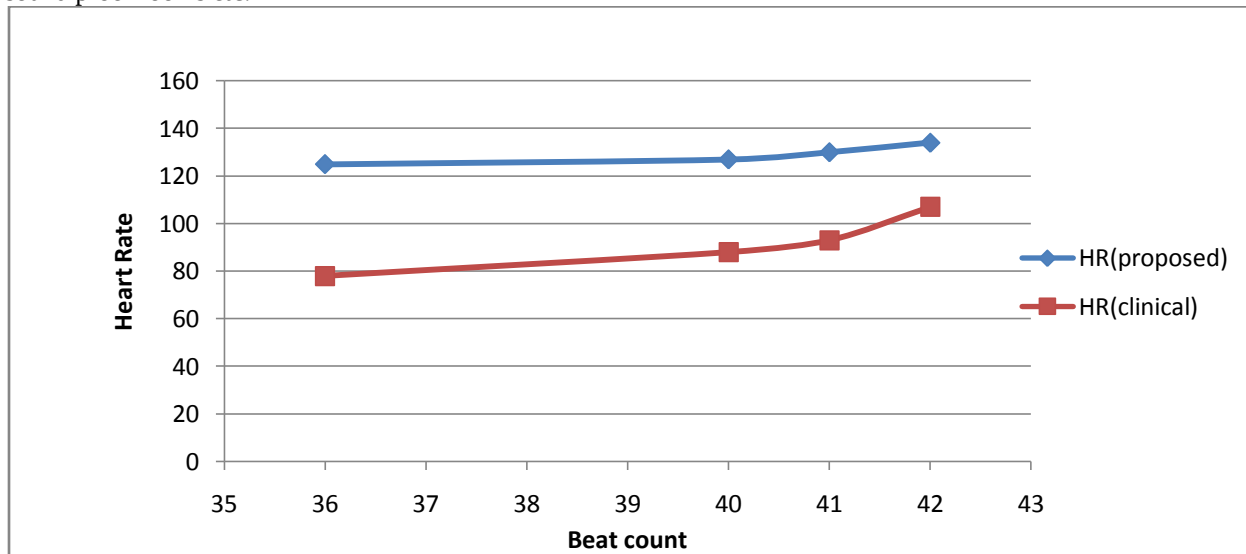


Figure 14 Comparison of Heart rate (proposed) to Heart Rate (clinical) [13]

REFERENCES

- [1] Meda Sai Kheerthana, Manjunath A.E, "A Survey on Wearable ECG Monitoring System Using Wireless Transmission of Data", *International Journal of Advanced Research in Computer and Communication Engineering*, Vol.4, Issue 7, July 2015.
- [2] Mesleh A, Skopin D, Baglikov S, Quteishat A, "Heart rate extraction from vowel speech signals", *Journal of Computer Science and Technology*, Volume 27, Issue 6, pp 1243–1251, 2012.
- [3] M.K Islam, A.N.M.M Haque, G.Tangim, T. Ahammad and M.R.H Khondokar, "Study and analysis of ECG signal using MATLAB and labview effective tools", *International journal of computer and electrical engineering*, Vol. 4, No. 3, june 2012.
- [4] J. Hada and Y. Takeuchi, "Evaluation of mental stress by voice analysis during simulated landing", *The Japanese Journal of Ergonomics*, Vol.44, No. 3 P171-174 (2008).

- [5] B. Schuller, F. Friedmann and Florian Eyben, "The Munich Biovoice Corpus: Effects of Physical Exercising, Heart Rate, and Skin Conductance on Human Speech Production", Proceedings of the Ninth International Conference on Language Resources and Evaluation, LREC (2014).
- [6] B. Schuller, F. Friedmann and Florian Eyben, "Automatic Recognition of Physiological Parameters in the Human Voice: Heart Rate and Skin Conductance", Proceedings of IEEE International Conference on Acoustics, Speech and Signal Processing, (2013).
- [7] Itkara F, "Minimum Prediction Residual Principle Applied to Speech Recognition", IEEE Transaction on acoustics, speech and signal processing, AS23(1):67 – 72, March 1975.
- [8] Skopin D, Baglikov S, " Heartbeat feature extraction from vowel speech signal using 2D spectrum representation", proceedings of IEEE International Conference Information Technology (2009).
- [9] A.P. James, "Heart rate monitoring using human speech spectral features", Human Centric computing and information sciences, (2015).
- [10] Mesleh A, Skopin D, Baglikov S, Quteishat A, " Heart rate extraction from vowel speech signals", Journal of Computer Science and Technology, Volume 27, Issue 6, pp 1243–1251, November 2012.
- [11] R.J Elliot, L. Aggoun and J.B, "Moore Hidden Markov Models: Estimation and control", Springer xii, 361 p. :ill. ; 25 cm (1995).
- [12] Alampreet Chahal, Silky Baghla, Gautam Kaushal and Dr. Rajbir kaur, "Heart Rate Monitoring using human speech featur extraction", International journal of engineering technology science and research, Vol 4, Issue 5, May 2017
- [13] Alampreet Chahal, Silki Baghla, Gautam Kaushal and Dr. Rajbir Kaur, "A simple approach for heart rate monitoring using Adaptive filters", Vol.4, Issue 6. June 2017.
- [14] Meraoumia A, Chitroub S, Bouridane A (2012) Multimodal biometric person recognition system based on fingerprint& finger-knuckle-print using correlation filter classifier. In: IEEE International Conference on Communications(ICC), 2012, pp 820–824
- [15] Breiman L , Bagging predictors. Mach Learn 24(2):123–140(1996)