

We are IntechOpen, the world's leading publisher of Open Access books Built by scientists, for scientists

6,900

Open access books available

186,000

International authors and editors

200M

Downloads

Our authors are among the

154

Countries delivered to

TOP 1%

most cited scientists

12.2%

Contributors from top 500 universities



WEB OF SCIENCE™

Selection of our books indexed in the Book Citation Index
in Web of Science™ Core Collection (BKCI)

Interested in publishing with us?
Contact book.department@intechopen.com

Numbers displayed above are based on latest data collected.
For more information visit www.intechopen.com



Silicon Technologies for Speaker Independent Speech Processing and Recognition Systems in Noisy Environments

Karthikeyan Natarajan¹, Dr.Mala John, Arun Selvaraj²
*Madras Institute of Technology, Anna University
 India*

1. Introduction

As the speaker independent speech recognition problem itself is highly computation intensive, the external environment adds to recognition complexity. As per Moore's law, doubling of number of transistors in a chip per year lead to the integration of various architectures in high density chips which lead to the implementation of high complex mixed signal speech systems in FPGA and ASIC technologies. Though several software based speech recognition systems are developed over the years, speech system implementations are yet to unleash the capabilities of silicon technologies. Direct mapped, completely hardware based systems will be highly energy efficient and less flexible but processor based implementation will be less energy efficient and flexible. Software based recognition systems fail to meet the latency requirements of the real time conditions whereas a completely hardware based recognition systems are power intensive. Hence in this case study, a hardware software based co-design is considered for the speech recognition implementation. Sequential algorithms which have been developed need to be modified to suit the parallel hardware systems. Hardware and software based co-design of the isolated word recognition problem will be applicable for low power systems like an AI based robotic system which could use a fixed point arithmetic and hence algorithmic optimizations needed to be considered to suit the actual hardware. Isolated word recognition problem can be split into three stages namely speech analysis, robust processing and final recognition stage. This hardware based speech recognition system is characterized for power and computation efficiency with the following parameters namely vocabulary size, robust speech recognition, speech variability, power and fixed point inefficiencies. This hardware system uses 50Mbps (Max 100Mhz) / 50Mhz NIOS 2 processor with WM8731 audio codec, DRAM controller, I2C controller, Avalon Bus bridge controller, ASIP matrix processor and parallel log Viterbi based hardware module implemented in ALTERA FPGA.

This chapter provides an Introduction to Hidden Markov model based speech Recognition. Relative merits and demerits of conventional Filter bank based feature extraction algorithm via windowed Fourier transform method is compared with a parallel linear predictive coding based CMOS implementation. Detailed description of the HMM based speech

¹ Author is currently working in IBM India Systems and Technology Engineering Labs

² Author is currently working in Wipro Technologies, Chennai.

recognition SOC chip is explained in this section. The robust processing step which involves the removal of external unintended noise component from speech signal and the novel Application specific matrix processor for noise removal based on signal subspace based Frobenious norm constrained algorithm are further discussed. This ASIP matrix solver consists of Singular value decomposition unit, QR decomposition unit, matrix bi-diagonalization unit, Levinson-Durbin Toeplitz matrix solver, fast matrix transposition unit based on efficient address generation module. Discussion on word recognition implementation as a parallel 32 bit fixed point 32 state univariate Hidden Markov model based system in ALTERA FPGA is carried out in the final section of this chapter.

1.1 Introduction to HMM based speech recognition system

Speech recognition can be classified into three categories namely Isolated, Connected and Continuous speech recognition systems. In an isolated word recognition system, each word is assumed to be surrounded by silence or background noise. This means that both sides of a word must have no speech input, making definite word boundaries easy to construct. This kind of recognition is mainly used in applications where only a specific digit or a word needs to be identified. Implementation of Isolated word recognition doesn't require any language information and it uses the minimum information about the source speech and has the low recognition accuracy for very large vocabulary. Connected speech (or more correctly 'connected utterances') recognition is similar to isolated word Recognition, but it allows several words/digits to be spoken together with minimal silence period between them. Longer phrases or utterances are therefore possible to be recognized. Continuous speech recognition is method for recognizing spontaneous speech. The system is able to recognize a sequence of connected words, which are not separated by pauses, in a sentence. This mode requires much more computation time and memory, and it is more difficult to operate when compared to isolated word recognition. A speaker-dependent system is a system that recognizes a specific speaker's speech while speaker-independent systems can be used to detect speech by any unspecified speaker. Currently speaker independent systems are modeled using Gaussian Mixture based quantizers which have high recognition accuracy. For speaker independent speech recognition system the training data must be exhaustive, which should incorporate all kinds of speaker variations. It is clear that the smaller the vocabulary size, the higher the recognition accuracy. In an isolated digit recognition system we can achieve higher accuracy by storing finer models of the digits. Further if the vocabulary size is increased there is significant reduction in the computational performance of the system. The training data needs to be generated from the field or the environment where we are planning to implement it.

Isolated Word Recognition problem can be divided into two parts, namely - Front End Processing and Pattern recognition. Typically, the front-end building block includes two modules, data acquisition and feature extraction. In our system we have also implemented the end-point detection and speech enhancement module to make the speech signal more adaptive and robust to the noise. The first stage in any Speech Recognition system is modeling the input speech signal based on certain objective parameters also called the Front End Parameters. Modeling of the input speech signal involves three basic operations spectral modeling, Feature extraction, and parametric transformation (Figure 1). Spectral shaping is the process of converting the speech signal from analog to digital and emphasizing important frequency components in the signal. Noise suppression and speech enhancement module can be added to the Front end processing module which will improve the recognition accuracy.

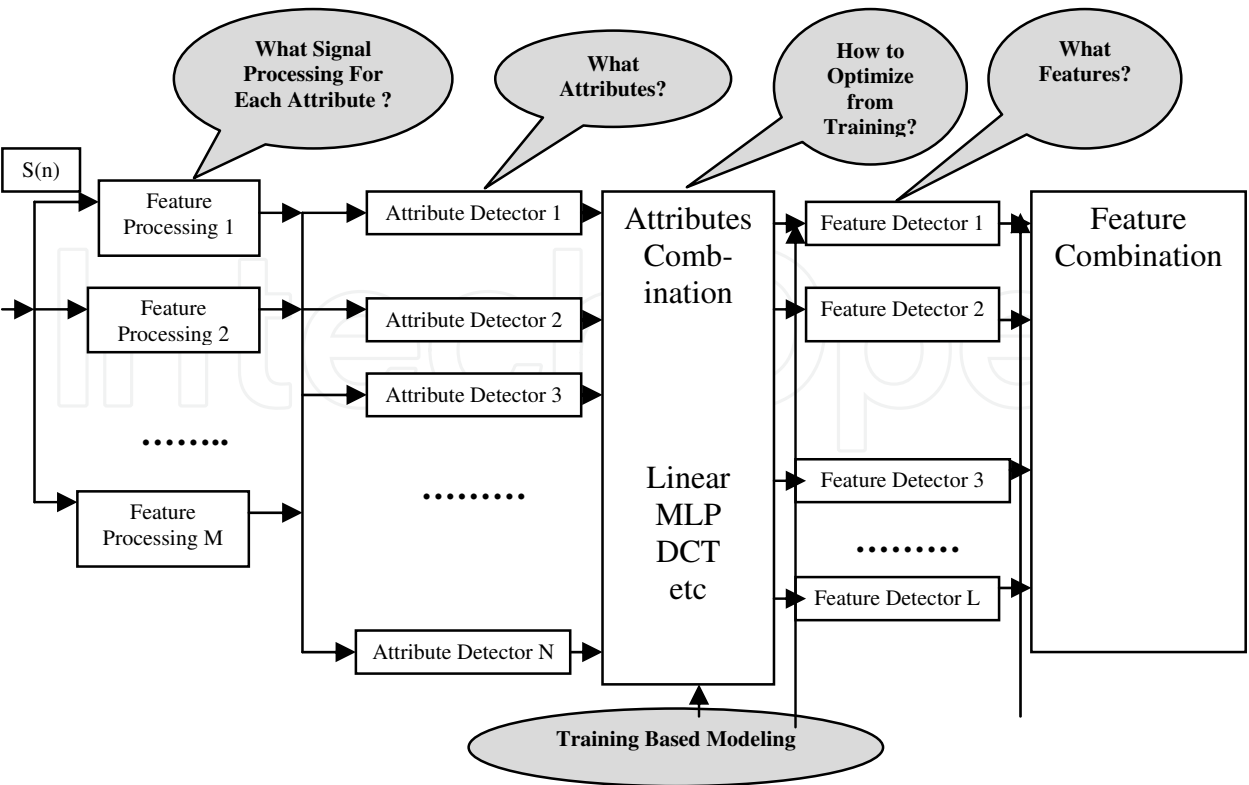


Fig. 1. Components of a speech recognition system

Two major kinds of front end Processing methods are Linear Predictive Coding and Mel Frequency Cepstral Co-efficient. The basic idea behind the linear predictive coding (LPC) analysis is that a speech sample can be approximated as a linear combination of past speech samples. By minimizing the sum of the squared differences (over a finite interval) between the actual speech samples and the linearly predicted ones, a unique set of predictor coefficients is determined. Speech is modeled as the output of linear, time-varying system excited by either quasi-periodic pulses (during voiced speech), or random noise (during unvoiced speech). The linear prediction method provides a robust, reliable, and accurate method for estimating the parameters that characterize the linear time-varying system representing vocal tract. In linear prediction (LP) the signal $s(n)$ is modeled as a linear combination of the previous samples:

$$s(n) = \overrightarrow{S}(n + e(n) = \sum_{i=1}^{NLP} a_{LP}(i)s(n - i) + e(n)$$

(1)

$a_{LP}(i)$ are the coefficients that need to be decided, N_{LP} is the order of the predictor, i.e. the number of coefficients in the model, and $e(n)$ is the model error, the residual. There exists several methods for calculating the coefficients. The coefficients of the model that approximates the signal within the analysis window (the frame) may be used as features, but usually further processing is applied. Higher the order of the LP Filters used, better will be the model prediction of the signal. A lower order model, on the other hand, captures the trend of the signal, ideally the formants. This gives a smoothened spectrum. The LP coefficients give uniform weighting to the whole spectrum, which is not consistent with the

human auditory system. For voiced regions of speech all pole model of LPC provides a good approximation to the vocal tract spectral envelope. During unvoiced and nasalized regions of speech the LPC model is less effective than voiced region. The computation involved in LPC processing is considerably less than cepstral analysis. Thus the importance of method lies in ability to provide accurate estimates of speech parameters, and in its relative speed. The features derived using cepstral analysis outperforms those that do not use it and that filter bank methods outperform LP methods. Best performance was achieved using MFCCs with Filter bank processing. Even though the CPU computations and memory accesses for MFCC are more, they are less speaker dependent and more speaker Independent. In our implementation we are using Short Time Fourier Transform based MFCC Feature Extraction Method for Front End Processing(Figure 2).

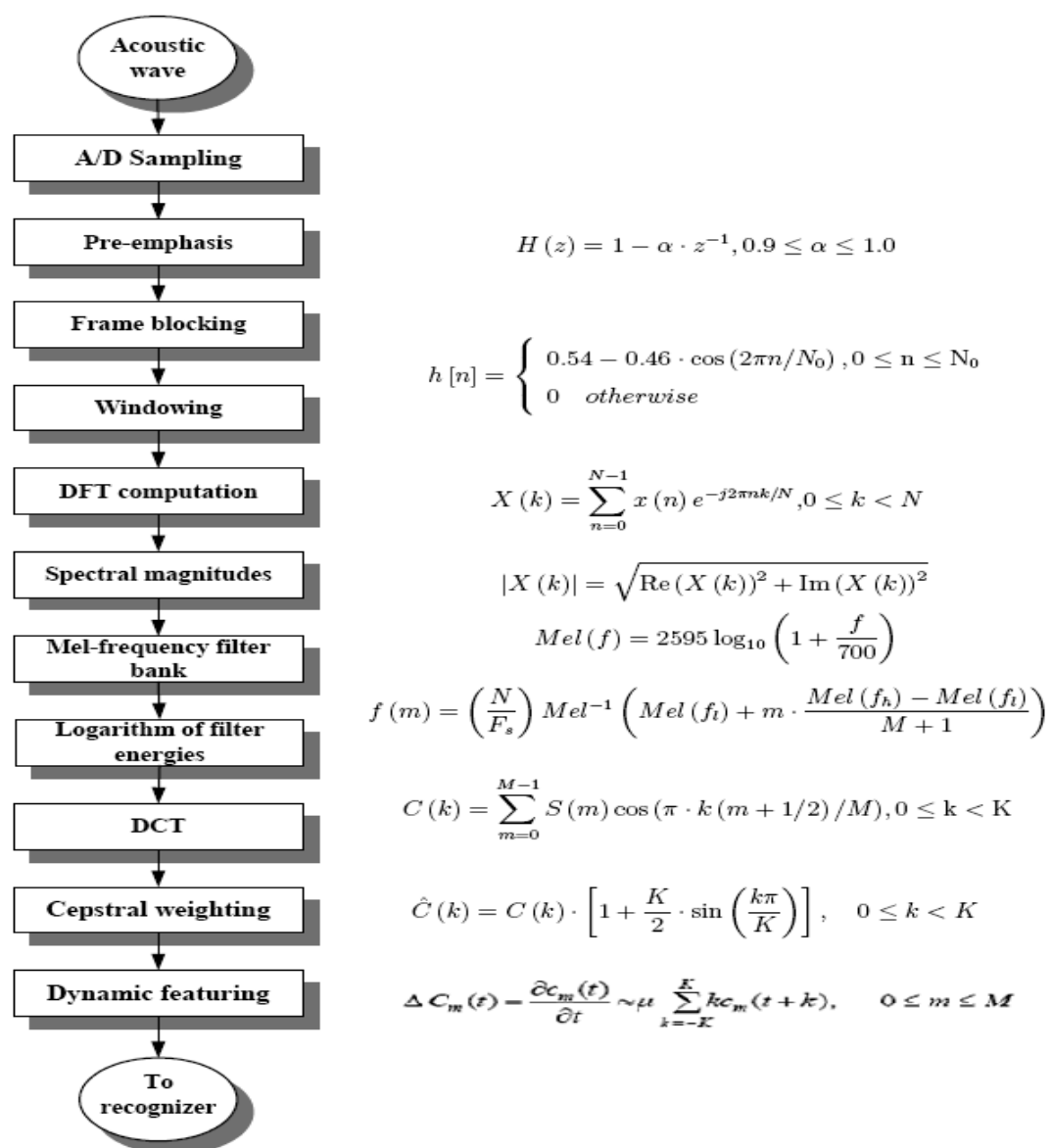


Fig. 2. Flow for Front End processing with feature extraction

We have found that with the hamming window of length 256 the signal can be represented efficiently with consideration to the hardware computational requirements of implementing

a FFT routine. After windowing the speech signal, Discrete Fourier Transform (DFT) is used to transfer these time-domain samples into frequency-domain ones. Direct computation of the DFT requires N^2 operations, assuming that the trigonometric functions have been pre-computed. Meanwhile, the FFT algorithm only requires on the order of $N \log_2 N$ operations, so it is widely used for speech processing to transfer speech data from time domain to frequency domain.

$$X(k) = \sum_{n=0}^{N-1} x(n) e^{-j2\pi nk/N} \quad 0 \leq k < N \quad (2)$$

The Spectral magnitudes were obtained by computing the absolute values of the FFT real and imaginary outputs. The square root is a monotonically increasing function and can be ignored if only the relative sizes of the magnitudes are of interest (ignoring the increased dynamic range).

$$|X(k)| = \sqrt{\text{Re}(X(k))^2 + \text{Im}(X(k))^2} \quad (3)$$

The computation still requires two real multiplications and consumes a lot of latency. A well-known approximation to the absolute value function is given as.

$$|A_{\text{re}} + jA_{\text{im}}| \approx |A_{\text{re}}| + |A_{\text{im}}| \quad (4)$$

A less frequently used approximation is only slightly more complex to implement but offers far better performance (refer table 1).

$$|A_{\text{re}} + jA_{\text{im}}| \approx \max(|A_{\text{re}}|, |A_{\text{im}}|) + \frac{1}{2} \min(|A_{\text{re}}|, |A_{\text{im}}|) \quad (5)$$

The above approximation was considered for the computation of spectral magnitude of the FFT outputs and their spectral magnitudes are taken. Human auditory system is nonlinear in amplitude as well as in frequency. We have taken logarithm to emulate amplitude nonlinearity and Mel filter banks to incorporate frequency nonlinearity. We have used 27 Mel triangular filter banks with 102 coefficients evenly spaced in Mel domain and the cepstral vectors are extracted based on the following equation 6 (refer Figure 3).

$$f(k) = (N / Fs) * \text{Mel}^{-1}(\text{Mel}(F_{\text{low}}) + k * \frac{(\text{Mel}(F_{\text{high}}) - \text{Mel}(F_{\text{low}}))}{M + 1}) \quad (6)$$

$$\text{Mel}(f) = 2595 * \log_{10}(1 + \frac{f}{700}) \quad (7)$$

$$\text{Mel}^{-1}(f) = 700 * (10^{\frac{f}{2595}} - 1) \quad (8)$$

$$H_m(k) = \begin{cases} 0 & \text{if } k < f(m-1) \\ \frac{(k - f(m-1))}{f(m) - f(m-1)} & \text{if } f(m-1) \leq k \leq f(m) \\ \frac{(f(m+1) - k)}{f(m+1) - f(m)} & \text{if } f(m) \leq k \leq f(m+1) \\ 0 & \text{if } k > f(m+1) \end{cases} \tag{9}$$

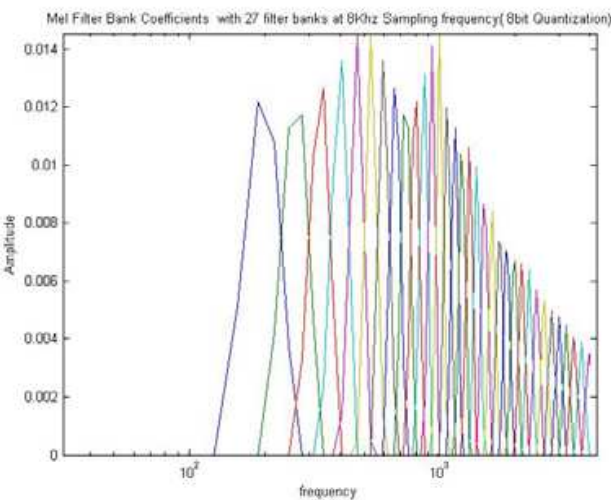


Fig. 3. Mel Filter Bank

The inverse DFT is performed on the output of the filter bank. Since the log power spectrum is symmetric and real, the inverse DFT is reduced to discrete cosine transformation (DCT). This transformation decorrelates features, which leads to using diagonal covariance matrices instead of full covariance matrices while modeling the feature coefficients by linear combinations of Gaussian functions. Therefore complexity and computational cost can be reduced. This is especially useful for speech recognition systems. Since DCT gathers most of the information in the signal to its lower order coefficients, by discarding the higher order coefficients, significant reduction in computational cost can be achieved. Typically the number of coefficients, K , for recognition ranges between 8 and 13. The equation is as following

Sensitivity of the lower order cepstral coefficients to overall slope and a higher order coefficient to noise has necessitated weighing of the cepstral coefficients by a tapered window to minimize these sensitivities. We have used weighing by a band pass filter of the form.

Temporal cepstral derivatives are an improved feature vector for forming the speech frames. They can be used with the cepstral derivative in case the cepstral Coefficients do not give acceptable recognition accuracy. Cepstral representations provide good approximations to the local spectral Properties. Derivatives of cepstral coefficients can be used to describe the dynamic movement of spectrum. In practical applications, the following approximation is used,

$$\Delta C_m(t) = \frac{\partial C_m(t)}{\partial t} \approx \left\{ \mu * \sum_{k=-K}^K k * C_m(t+k) \right\} \quad 0 \leq m \leq M \tag{10}$$

Where μ is a normalization factor.

Typical feature vector: (Figure 4):

$[E(t) \ c_1(t) \ c_2(t) \dots \ c_M(t), \ E(t) , \ \Delta E(t) , \ \Delta c_1(t) \ \Delta c_2(t) \dots \ \Delta c_M(t-1) \ \Delta \Delta c_1(t) \ \Delta \Delta c_2(t) \dots \ \Delta \Delta c_M(t-1)]^T$

Feature vector consists of both static part and the Dynamic part of the speech signal.

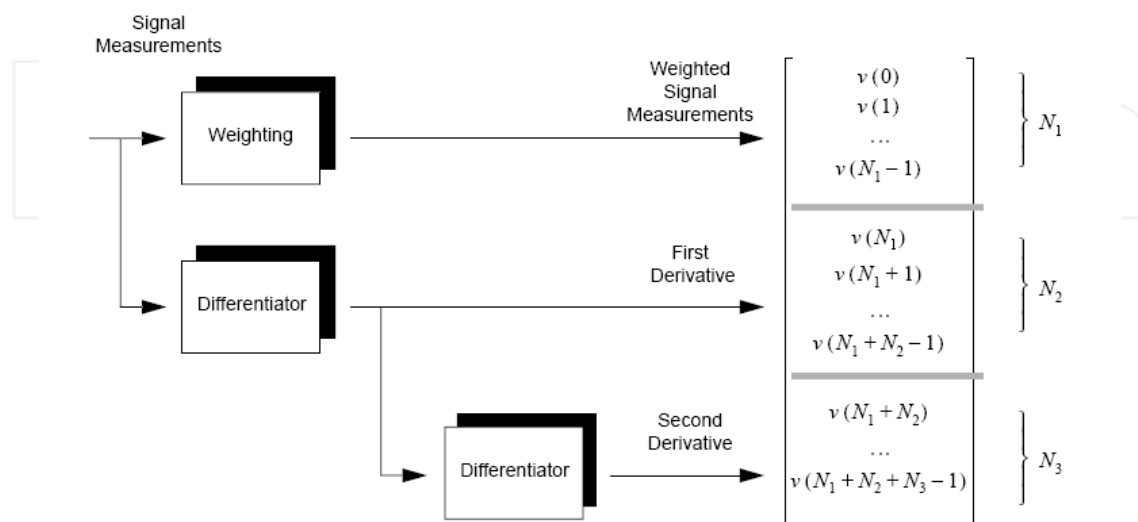


Fig. 4. Representation of Delta and Delta- Delta parameters

2. The hidden Markov models

2.1 The three basic problems of HMM

1. Given the observation sequence $O = (o_1 \ o_2 \ \dots \ o_T)$, and a model $\lambda = (A, B, \pi)$, how do we efficiently compute $p(o | \lambda)$, the probability of the observation sequence, given the model. This is the "evaluation problem". Using the forward and backward procedure provides solution.
2. Given the observation sequence $O = (o_1 \ o_2 \ \dots \ o_T)$, and the model λ , how do we choose a corresponding state sequence $q = (q_1, q_2, \dots, q_T)$ that is optimal in some sense (i.e. best explains the observation). The Viterbi algorithm provides a solution to find the optimal path.
3. How do we adjust the model parameters $\lambda = (A, B, \pi)$ to maximize $p(O | \lambda)$. This is by far the most difficult problem of HMM. We choose $\lambda = (A, B, \pi)$ in such a way that its likelihood, $p(O | \lambda)$, is locally maximized using an iterative procedure like Baum-Welch method (L. Rabiner 1993).

The base speech recognizer works only with noiseless HMM states and the matrix processor is used as pre conditioning block to generate the noiseless HMM models from the noisy speech vectors. There are three kinds of Hidden Markov models described in the literature namely Discrete HMM, Continuous HMM and Semi-continuous HMM (Vaseghi) in which a Continuous HMM model is used to model the HMM states. A HMM model is characterized by the no of states N , no of distinct observation symbols M , the transition probability matrix A , the initial probability matrix Π , the output observation probability for a feature x_1 in state I , $b_i(x_1)$.

2.2 The log-viterbi algorithm:

1) Initialization:

$$\delta^{(\log)}_1(i) = \log a_{1i} + \log b_i(x_1) \quad (11)$$

$$\psi_1(i) = 0 \quad (12)$$

2) Recursion:

$$\delta^{(\log)}_{t+1}(j) = \max_{i=2}^{N-1} (\delta^{(\log)}_t(i) + \log a_{ij}) + \log b_j(x_{t+1}) \quad (13)$$

$$\psi_{t+1}(i) = \arg \max_{i=2}^{N-1} (\delta^{(\log)}_t(i) + \log a_{ij}) \quad (14)$$

3) Termination:

$$\log(P(O/\lambda)) = \max_{i=2}^{N-1} (\delta^{(\log)}_T(i) + \log a_{iN}) \quad (15)$$

$$q^T = \arg \max_{i=2}^{N-1} (\delta^{(\log)}_T(i) + \log a_{iN}) \quad (16)$$

4) Backtracking:

$$q_t = \psi_{t+1}(q_{t+1}) \text{ for } t=T-1 \text{ to } 1 \quad (17)$$

The probability of observation vectors, $p(O|\lambda)$ has to be maximized for different model parameter values which corresponds to HMM models for different words. The implementation of the log likely computation can be done in an efficient way using the Forward and backward procedures as described in (Karthikeyan - ASICON 2007). Since the direct implementation of Viterbi algorithm results in underflow due to very low probability values are multiplied recursively over the speech frame window, logarithmic Viterbi algorithm is implemented which is different from methods given in (Karthikeyan - ASICON 2007). Since the direct implementation of Forward, Backward as well as the Viterbi algorithm results in underflow, we took logarithm on both sides and we have implemented logarithmic versions of the above algorithm. Since the Forward algorithm uses summation which is being replaced by the following conversion in the modified forward algorithm. We have used the modified forward algorithm, backward algorithm as well as viterbi algorithm which is different from the methods given in [6].

2.3 The Baum Welch re-estimation procedure

The third, and by far the most difficult, problem of HMMs is to determine a method to adjust the model parameters (A, B, π) to maximize the probability of the observation sequence given the model. There is no known way to analytically solve for the model, which maximizes the probability of the observation sequence. In fact, given any finite observation sequence as training data, there is no optimal way of estimating the model parameters. We can, however, choose $\lambda = (A, B, \pi)$ such that $P(O|\lambda)$ is locally maximized using an iterative procedure such as the Baum-Welch method. To describe the procedure for re-estimation

(iterative update and improvement) of HMM parameters, we first define $\xi_t(i,j)$, the probability of being in state S_i at time t , and state S_j , at time $t+1$, given the model and the observation sequence.

In order to use either ML or MAP classification rules, we need to create a model of the probability $p(o_j)$ for each of the different possible classes. The PDF can be modeled using a Gaussian distribution. We can create a Gaussian model by just finding the sample Mean, and the sample covariance matrix U_i .

$$\begin{aligned}\mu_i &= \frac{1}{N} \sum_{n=1}^N o_n \\ U_i &= \frac{1}{N-1} \sum_{n=1}^N (o_n - \mu_i)'(o_n - \mu_i)\end{aligned}\quad (18)$$

$$N(o; \mu, U) = \frac{1}{\sqrt{(2\pi)^p |U|}} \exp \left(-\frac{1}{2} (o - \mu) U^{-1} (o - \mu)' \right) \quad (19)$$

Probability of being in state S_i at time t , and state S_j at time $t+1$, given the model and the observation sequence, i.e.

$$\xi_t(i, j) = P(q_t = S_i, q_{t+1} = S_j | O, \lambda). \quad (20)$$

2.4 Covariance selection in speech recognition

The covariance Matrix used in model based speech Recognition problem which uses N state Univariate Gaussian HMM modeling with M dimensional features can be considered in the following ways. The following 39 dimensional feature vectors are considered for designing the continuous HMM based speech recognizer.

$[c_1(t) \ c_2(t) \dots c_M(t), \ \Delta c_1(t) \ \Delta c_2(t) \dots \Delta \Delta c_M(t-1) \ \Delta \Delta c_1(t) \ \Delta \Delta c_2(t) \dots \Delta \Delta c_M(t-1), \ E(t), \ \Delta E(t)]^T$

Where $\Delta C_m(t)$, $\Delta \Delta C_m(t)$ can be represented as below.

$$\Delta C_m(t) = \frac{\partial C_m(t)}{\partial t} \approx \left\{ \mu^* \sum_{k=-K}^K k * C_m(t+k) \right\} 0 \leq m \leq M \quad (21)$$

$$\Delta \Delta C_m(t) = \frac{\partial \Delta C_m(t)}{\partial t} \approx \left\{ \mu^* \sum_{k=-K}^K k * \Delta C_m(t+k) \right\} 0 \leq m \leq M \quad (22)$$

- i. Complete covariance matrix (distance measure: Mahalanobis distance measure)
Complete covariance matrix when considered results in very high implementation complexity and cannot be easily achieved with the existing hardware resources (S.Yozhizawa – 2006).
- ii. The second method is to do covariance parameter tying (Pihl – 1996). In this a method a common parameter is considered for all the states and other statistical characteristics are considered different. For instance, use of common covariance matrix for all the

- clusters obtained during GMM block quantization and considering mean, no of observation output different for each state.
- iii. A still modified version of the above mentioned method having block covariance matrices instead of having complete covariance matrix can be considered for implementation which has complex implementation in hardware. Our assumption that the covariance are block diagonal is valid since the use of an orthogonal transform like DCT decorrelates the cepstral vectors. The correlation only exists between the time difference cepstral vectors, delta cepstral vectors, and the delta - delta cepstral vectors. So we can construct the covariance matrix as three element block diagonal matrix for which the inverse matrix can be easily found using Singular value decomposition.
 - iv. The last method is to consider the covariance matrix to be diagonal which yields the simplest hardware architecture. The inverse diagonal values are stored in memory locations and only multiply operations are performed and this method is computationally less intensive. Present hardware based recognizers implement this method due to its low complexity. But this method will produce significant errors and degrade the recognition performance of the system as it is doesn't efficiently represent the correlation introduced by the Vector quantizer. Earlier proposed implementations were based on this method only (Karthikeyan – ASICON 2007). Where $E(.,.)$ represents the statistical Expectation operation on the cepstral vectors.

$$R = \begin{pmatrix} \begin{pmatrix} E(c_1, c_1) & E(c_1, \Delta c_1) & E(c_1, \Delta\Delta c_1) \\ E(\Delta c_1, c_1) & E(\Delta c_1, \Delta c_1) & E(\Delta c_1, \Delta\Delta c_1) \\ E(\Delta\Delta c_1, c_1) & E(\Delta\Delta c_1, \Delta c_1) & E(\Delta\Delta c_1, \Delta\Delta c_1) \end{pmatrix} & 0 & 0 \\ 0 & \begin{pmatrix} E(c_2, c_2) & E(c_2, \Delta c_2) & E(c_2, \Delta\Delta c_2) \\ E(\Delta c_2, c_2) & E(\Delta c_2, \Delta c_2) & E(\Delta c_2, \Delta\Delta c_2) \\ E(\Delta\Delta c_2, c_2) & E(\Delta\Delta c_2, \Delta c_2) & E(\Delta\Delta c_2, \Delta\Delta c_2) \end{pmatrix} & 0 \\ 0 & 0 & \begin{pmatrix} E(c_3, c_3) & E(c_3, \Delta c_3) & E(c_3, \Delta\Delta c_3) \\ E(\Delta c_3, c_3) & E(\Delta c_3, \Delta c_3) & E(\Delta c_3, \Delta\Delta c_3) \\ E(\Delta\Delta c_3, c_3) & E(\Delta\Delta c_3, \Delta c_3) & E(\Delta\Delta c_3, \Delta\Delta c_3) \end{pmatrix} \end{pmatrix} \quad (23)$$

Earlier we have discussed the influence of the covariance selection on the performance of the recognition and it directly influences the word error rate. Earlier implementation considers completely diagonal co variances which cause drastic errors as we have introduced correlation into the feature vectors through vector quantization as well as dynamic feature vector set. Hence we can consider the feature vectors to be correlated only among the two dynamic features set delta and delta delta feature vectors the static features. Hence we can assume the correlation matrix to be block diagonal and hence the inverse of such a matrix can be easily obtained by linear equation solvers. Computation of the Singular Value Decomposition of a matrix A can be accelerated by the parallel two sided jacobi method with some pre-processing steps which would concentrate the Frobenius norm near the diagonal. Such approach would help noise reduction is great way as the noise sub-space is computed with Frobenius norm constraints. Such a concentration should hopefully lead to fewer outer parallel iteration steps needed for the convergence of the entire algorithm. However the gain in speed as measured by total parallel execution time depends decisively on how efficient is the implementation of the distributed QR and LQ factorizations on a given parallel architecture.

3. Speech recognition architecture

3.1 NIOS embedded processor based system design

NIOS 2 is a soft Processor which can be realized in any of the Altera's FPGA Development kits. It is based on a 32-bit RISC architecture and is a natural choice in projects where CPU performance is essential. The NIOS processor can be run at different frequencies based on which the Computational capability of the processor can be chosen. Nios Processor is available in three different speed grades and can be extended with additional coprocessors, instruction sets, and so forth. By doing so, it is possible to develop a large part of the system on an ordinary PC running Windows or any variant of UNIX. By simulating IP cores (firmware modules) as software objects, a system can be developed to an advanced state before it needs to be tested on the actual target. Another benefit of this approach is that it allows concurrent development of multiple projects on a single target. The NIOS processor is a 32-bit Harvard Reduced Instruction Set Computer (RISC) architecture optimized for implementation in Altera FPGAs with separate 32-bit instruction and data buses running at full speed to execute programs and access data from both on-chip and external memory at the same time. Nios Processor has got 32 32bit general purpose registers and 16 32bit control registers, an Arithmetic Logic Unit (ALU), Exception Unit, Instruction cache and Data Cache, Hardware multiply and Hardware divide, a barrel shifter unit and 32 software interrupts. This flexibility allows the user to balance the required performance of the target application against the logic area cost of the soft processor. NIOS processor does not separate between data accesses to I/O and memory (i.e. it uses memory mapped I/O). All the system peripherals of Altera are connected through a system bus called Avalon. for sophisticated SOPC environment or the basic environment. The stack convention used in Nios processor starts from a higher memory location and grows downward to lower memory locations when items are pushed onto a stack with a function call. Items are popped off the stack the reverse order they were put on; item at the lowest memory location of the stack goes first and etc (NIOS 2006). Nios processor also supports reset, interrupt, user exception, and break and hardware exceptions. The processor will only react to interrupts if the Interrupt Enable (IE) bit in the Machine Status Register (MSR) is set to 1. On an interrupt the instruction in the execution stage will complete, one has to manually enable the interrupt enable bit for that particular device.

Writing software to control the NIOS processor must be done in C/C++ language. The NIOS tool has got gnu based have built in C/C++ compilers and debugger to generate the necessary machine code for the NIOS processor (Agarwal 2001). NIOS Processor supports word (32 bits), half-word (16 bits), and byte accesses to data memory. Data accesses must be aligned (i.e. word accesses must be on word boundaries, half-word on half-word boundaries), unless the processor is configured to support unaligned exceptions. All instruction accesses must be word aligned.

Avalon Bus system is a simple yet extremely powerful bus system which allows any no of Bus masters to be added simultaneously and offers excellent arbitration capabilities with wait cycles. It also supports a unique kind of hardware software interface called custom instruction which acts as a hardware mapped instruction to the NIOS processor (Avalon 2006). We can also accelerate the software function in NIOS processor through a technology called Custom instruction which is unique to NIOS based system. Nearly 256 custom instructions different cycle times can be integrated into the design to accelerate the underlying software. Compared to complete software performance, a system with hardware

acceleration improves 20X performance improvement. Our design utilizes this custom instruction feature (Avalon 2006). NIOS processor supports many software IP-cores such as Timer, Programmable counters, Ethernet controller, DRAM controller, Flash controller, User logic components, PLLs, Hardware Mutex, LCD controller etc. NIOS Timers can be used to compute the execution time of a software routine or used to produce trigger at regular intervals so as to signal some of the hardware peripherals. Hardware IP cores can be connected to the system in two different ways. The hardware component can be configured as Avalon Custom instruction component and the processor can access the hardware as though it being an instruction. NIOS processor supports four different kinds of custom instruction technology namely combinational; Multi cycle, Extended and Internal Register file based custom instruction. Custom instruction module can also be connected to the Avalon Bus so that one can connect some of the custom instruction signals to external signals not related to the processor signals. Hardware IP core can also be interfaced to the NIOS system through the Avalon slave or Master Interface. Avalon Slave devices can have interrupts and they request the service of the processor through the interrupts. These interrupts can be prioritized manually.

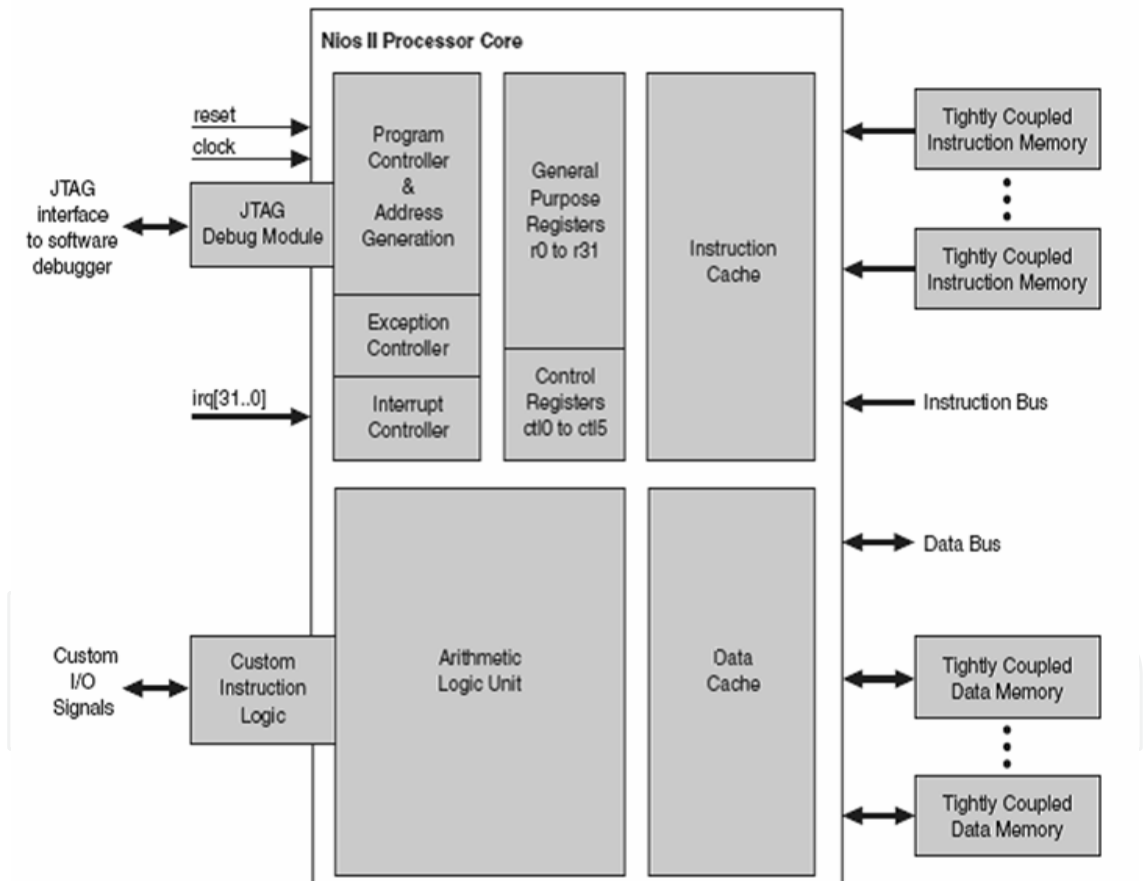


Fig. 5. NIOS Architecture

4. Design implementation using fixed point architecture

4.1 Finite word length effects

All DSP based designs strongly depend on the floating point to fixed point conversion stage as the DSP algorithm may not be implementable in floating point form. Fixed pint analysis

of the system is extremely important to understand the nonlinear nature of the quantization characteristics. This leads to certain constraints and assumptions on quantization errors: for example that the word-length of all signals is the same, that quantization is performed after multiplication, and that the word-length before quantization is much greater than that following Quantization (Meng 2004). Error signals, assumed to be uniformly distributed, with a white spectrum and uncorrelated, are added whenever a truncation occurs. This approximate model has served very well, since quantization error power is dramatically affected by word-length in a uniform word-length structure, decreasing at approximately 6dB per bit. This means that it is not necessary to have highly accurate models of quantization error power in order to predict the required signal width. In a multiple word-length system realization, the implementation error power may be adjusted much more finely, and so the resulting implementation tends to be more sensitive to errors in estimation. Signal-to-noise ratio (SNR), sometimes referred to as signal-to-quantization noise ratio (SQNR), is The ratio of the output power resulting from an infinite precision implementation to the fixed-point error power of a specific implementation defines the signal-to-noise ratio In order to predict the quantization effect of a particular word-length and scaling annotation, it is necessary to propagate the word-length values and scaling from the inputs of each atomic operation to the operation output (Haykin 1992). The precision of the output not only depends on the binary precision of the inputs, it also depends on the algorithm to be implemented. For example the fixed point implementation of complex FFT algorithm decreases 0.5 bit precision for each stage of computation (Baese 2005). So for large FFT lengths more bits of precision are lost. The Feature extraction stage was implemented in Nios Processor with fixed precession inputs. The following plots describe the fixed point characteristics of the algorithm.

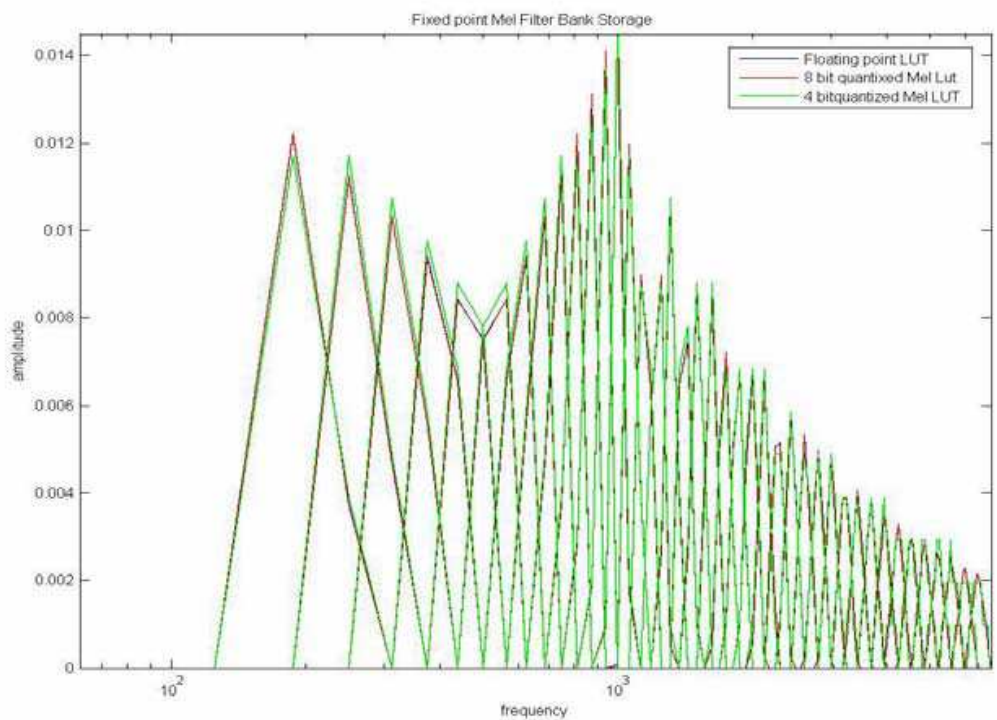


Fig. 6. Fixed point MFCC implementation

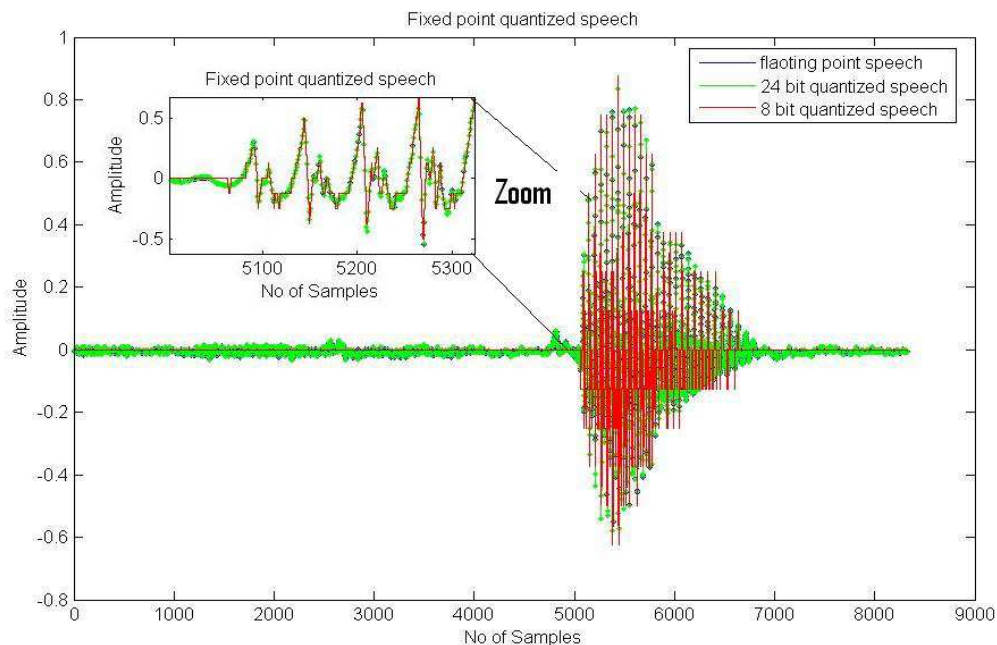


Fig. 8. Fixed point Speech input characteristics

4.2 Flexibility

The recognition system must be able to operate under a variety of conditions (Vaseghi 2004). The signal to noise ratio may vary significantly, the word may be stretched too long or too short, some of the states may be skipped, noise content may be high and we are forced to model the noise HMM and subtract it from the actual speech HMM (Hermus 2007). The receiver must incorporate enough programmable parameters to be reconfigurable to take best advantage of each situation (Press 1992). We applied signal subspace based noise reduction algorithm based on Singular Value decomposition to reduce the noise characteristics of the speech signal (Hemkumar 1991).

4.3 FCTSVD algorithm (Abut 2005):

1. Estimate the noise vectors W from silence periods in the observed speech signal.
2. Form the Hankel matrix H_y from the observed speech signal.
3. Compute SVD of H_y .
4. Initialize the order of retained singular values of H_x —
5. Let $S=S+1$ and reconstruct the estimated matrix of H_x , H_x^- using the first s eigen values.
6. Compute Frobenius Constrained norm metric and error is less than 0.0098 else goto 5.

4.4 Scaling:

As the number of frames in speech increases, the values of the variable the formal algorithm saturates whereas the log-Viterbi algorithm used in this hardware does not suffer this problem as the implementation involves only additions rather than multiplication.

4.5 Initial estimates of HMM parameters:

Uniform initial estimates were used for A and P_i Matrices. However B matrix cannot be initialized with random values as it has more influence in convergence of the Baum-Welch

algorithm. Since continuous hidden markov models are used, the initial estimates of B, Mean, and Variance are obtained using segmental K-means algorithm..

5. Project modules:

5.1 Main modules:

1. First module is concerned with signal analysis and feature extraction (FRONT END PROCESSING → SOFTWARE EXECUTED IN NIOS 2 PROCESSOR).
2. The next module generates the values of parameters required for comparing the test speech signal with the reference signal values .Training phase and this step should be robust since it directly determines the accuracy and the application where the system is to be deployed. (TRAINING – OFFLINE DONE IN MATLAB-refer Figure 13).
3. Maximum Likelihood based word recognition (PARALLEL HARDWARE).

5.2 Supporting hardware modules:

1. Audio Codec Configuration unit based on I2C controller (MPU 2 AUDIO CODEC CONFIGURATION).
2. Custom Avalon Master Module for Audio Codec data Retrieval with integrated SRAM memory controller.
3. Custom Speech controller for hardware recognition part with efficient mode management unit based on FSMs.
4. The speech Controller has got the following modules built in:
5. Viterbi based Speech Recognition unit with memory controllers for Model parameter RAMS.
 - a. Input Frame buffers for feature storage with memory controller for feature storage RAM.
 - b. Output Frame buffer for Model output storage
 - c. Efficient mode management unit to switch between various modes of operations using FSMs.
 - d. LED Display unit to finally display the results.
6. Custom Singular Value Decomposition unit.

Software Modules	Hardware Modules interfaced to Avalon	Custom Instructions
✓ Voice Activity Detection ✓ FFT based feature extraction ✓ Mel Filter Banks ✓ Cepstral Weighing ✓ Software Back-substitution	✓ Audio Serial 2 Parallel Module(Avalon Master) ❖ Input Frame buffers ❖ Output Frame Buffers ❖ Speech Recognition Mode controller top ❖ for SVD Speech Recognition module with Viterbi.	• Max comparator for magnitude computation in FFT

Table 2. Isolated Word Recognition System Hardware/ Software Partition

- Operating frequency of green = 18.432MHz
- ❖ Operating frequency of red= 12.5MHz

✓ **Operating frequency of 50 MHz**

Audio codec is configured via I2C interface. Two signals I2C_clock and I2C_data are used to configure the internal registers of the WM8731. The WM8731 is a WRITE-ONLY device; any requests to read are ignored. Device is configured by writing data to internal registers. The internal registers are configured by transferring data and address of the internal registers serially through I2C_data pin. Clock signal is applied to the I2C_clk pin. Clock signal can be generated in two modes of operation namely USB/Normal mode master clock (AUD_XCLK, from which AUD_BCLK is generated). USB mode must have a FIXED AUD_XCLK of 12MHz, which can be easily obtained from a PLL in the SOPC system. Normal mode requires AUD_XCLK clocks of either 12.288MHz (8kHz, 32kHz, 48kHz, 96kHz) or 11.2896MHz (44.1kHz, 88.2kHz). This implementation utilizes Normal mode of clock generation at 18.432MHz. Transfer is initiated by pulling MPU_DATA low while MPU_CLK is high. The data format of the configuration of a particular internal register has got 3-bytes.

- Byte 1: {ADDR[6..0],0} → ADDR[6..0] is DEVICE ADDRESS, which is ALWAYS 0x34
- Last bit is R/W bit, which is always 0 (write,) since WM8731 is write-only
- Byte 2: {REG[6..0],DATA[8]} → REG[6..0] is 7-bit register address, DATA[8] is MSB of 9-bit DATA
- Byte 3: DATA[7..0] → Lower 8 bits of 9-bit DATA
- MPU_DATA is driven low by the CODEC between bytes as confirmation

The following operations needed to be done to make the device operate in the intended mode of operation:

- Reset device: Write 0x00 to AUDIO_RESET
- Power up device: Write '0' to WM8731_POWER_DOWN_CTL.7 bit
- Disable Line IN -> Line OUT bypass, select MIC_IN (AUDIO_ANALOG_PATH_CTL Reg)
- Turn on MASTER mode: AUDIO_INTERFACE_FMT

5.3 How this hardware system works:

The steps involved in implementing this system consist of the following steps given below:

Step 1: Audio Codec is configured via CPU 2 I2C interface with the following specifications.

- ✓ **WM8731_POWER_DOWN_CTL** is used to power up the device.
- ✓ **WM8731_ANALOG_PATH_CTL** Register is set to 16'h08FD to enable the MIC in facility.
- ✓ **WM8731_SAMPLING_CTL** Register is set to 16'h100E to fix the audio codec in NORMAL MODE with ADC sampling frequency of 8 KHz. Codec operating frequency is 18.432MHz

Step2: The serial input bit stream is converted in parallel data using a custom Avalon Master interface and is stored in SRAM module. The storage of audio will be interrupted by a external user controlled switch to start the processing step.

Step3: This switch will induce an interrupt signal present in the speech recognition module (AVALON SLAVE CONFIGURED) to start the feature processing of NIOS processor.

Step4: In software the speech start and end points are detected, we perform windowing (Hamming Window).

Step5: We use short time Fourier analysis on the speech signal, since speech is a Quasi-periodic signal we need to use STFT. We have used a window of duration 30ms with an overlap of 10ms.

Step6: Evaluate the distance between the speech signals and do clustering using the Gaussian Mixture based Block quantizer based on Mahalanobis distance and clustering is performed.

Step7: The features are extracted and stored in the **INPUT FRAME BUFFER** of the Speech Recognition module.

Step8: Steps 1 to 6 will continue until the end of frame is detected by the hardware module.

Step9: Starting of speech recognition in hardware and finally the results are populated and displayed in LED. Each stage output is stored in **OUTPUT FRAME BUFFER** and final recognition is done.

5.5 Implementation of continuous hidden Markov model:

Our architecture concentrates on the three major issues Power, Memory access (Throughput) and vocabulary size. There is always a trade off existing between the operating frequency and the recognition vocabulary, word accuracy, noise suppression etc. This is a word HMM based architecture which uses continuous HMM for the implementation.

Two essential steps in the recognition algorithm are:

1. Output probability calculation.
2. Log VITERBI implementation.

→Output Probability calculation is the computationally intensive process as we need to do lots of multiplies and Add operations.

→Viterbi Algorithm is also implemented as a parallel processing block for faster recognition.

5.6 Hardware design:

Our architecture (Fig 11) concentrates on the three major issues Power, Memory access (Throughput) and vocabulary size. There is always a trade off existing between the operating frequency and the recognition vocabulary, word accuracy, noise suppression etc (Pihl 1996). This is a word HMM based architecture which uses continuous HMM for the implementation (Cho 2002).

Two essential steps in the recognition algorithm are:

1. Output probability calculation.
2. Log Viterbi implementation(as in fig 12).

5.7 Modes of operation:

We can operate the system in two modes:

1. Small vocabulary mode
2. Large vocabulary mode

5.8 Small vocabulary mode:

Startup Sequence:

1. *Reset* = 0, *sw0*= 0, *sw1*=0
2. Audio is stored in SRAM by custom master Avalon interface.
3. When the user presses switch0 the Avalon master stops storing samples.
4. Speech controller interrupts processor for features after *sw0* is pressed.
5. Processor starts processing the samples to extract features and once is complete activates done signal of Speech controller.

- 6. Load the models in the Model Rams.
- 7. Perform Viterbi and output probability computation (refer Figure 9 , Figure 10)
- 8. Store the results of word I in the Output Frame Buffers.
- 9. repeat the step until all word are done
- 10. Generate complete word signal
- 11. Compare the results and show the word spelt in LED.

5.9 Large vocabulary mode:

- 1. *Reset* = 0, *sw0*= 0, *sw1*=1
- 2. Audio is stored in SRAM by custom master Avalon interface.
- 3. Speech controller interrupts processor for features after the End of VAD is received.
- 4. Processor will load the models of 10 words from external memory to the INPUT FRAME BUFFERS OF SPEECH RECOGNITION MODULE.
- 5. Perform Viterbi and output probability computation.
- 6. Store the results of word I in the Output Frame Buffers.
- 7. Repeat the step until all word is done.
- 8. Generate complete word signal.
- 9. Compare the results and show the word spelt in LED.

Software algorithm in Floating point C		Average no of cycles to execute(ps)
Convolution		49484981.2
VAD		188143.2
Multiplication		1159.5
FFT		429496674.2
Twiddle Factors		44562.1
Magnitude	Actual (equation 2.8)	9115.8
	Approximation1(equation 2.9)	2545.4
	Approximation1(equation 2.10)	962.8
	Total computation	1224.1

Table 1. Software Execution Cycles of the Feature Extraction Algorithm

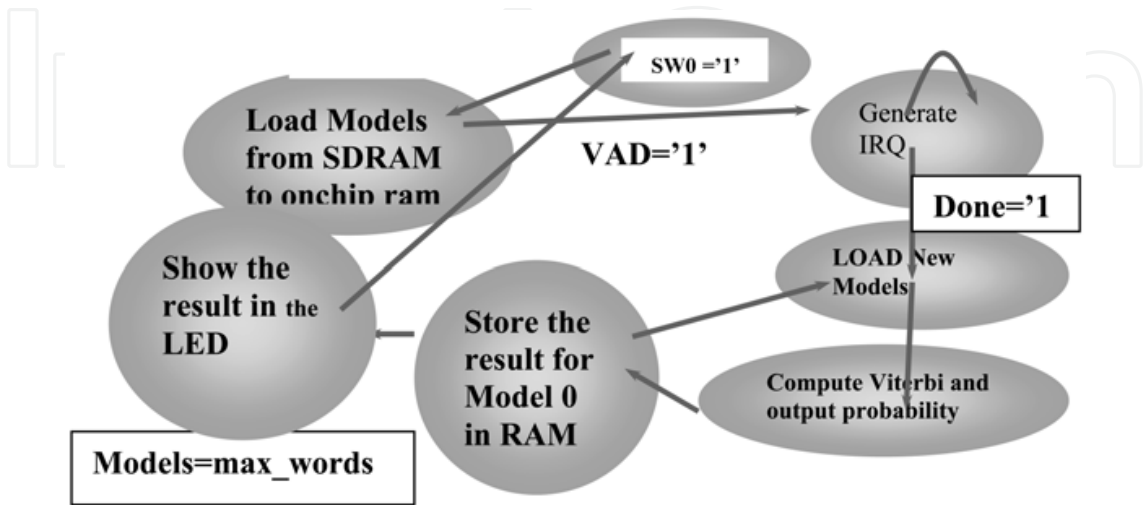


Fig. 9. FSM for main speech Recognition module.

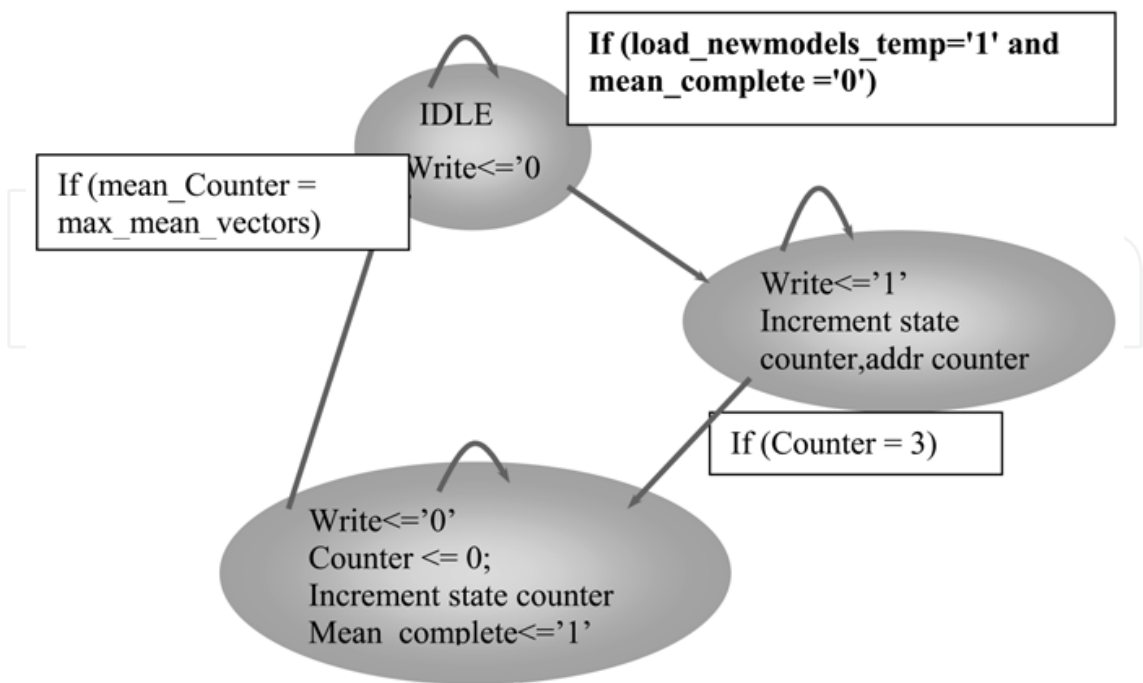


Fig. 10. FSM for memory access of (mean/Covariance Rom/transmit/initial) RAMs.

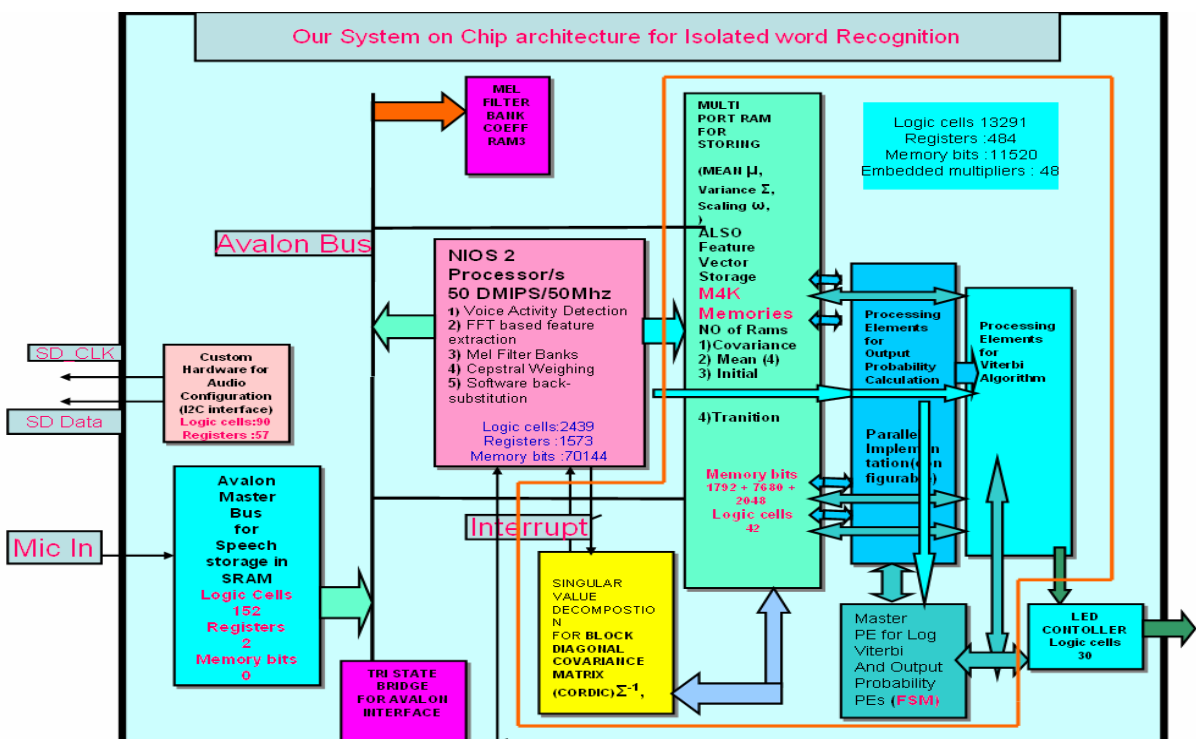


Fig. 11. The speech recognition hardware architecture

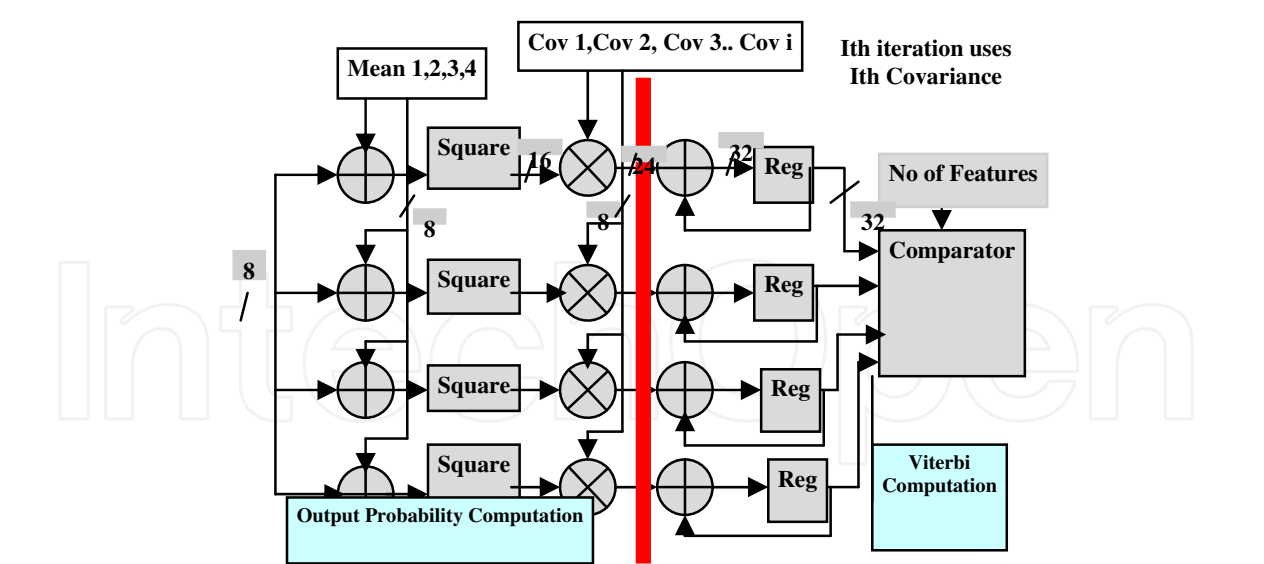


Figure 12. Pipelined-parallel Log likelihood unit

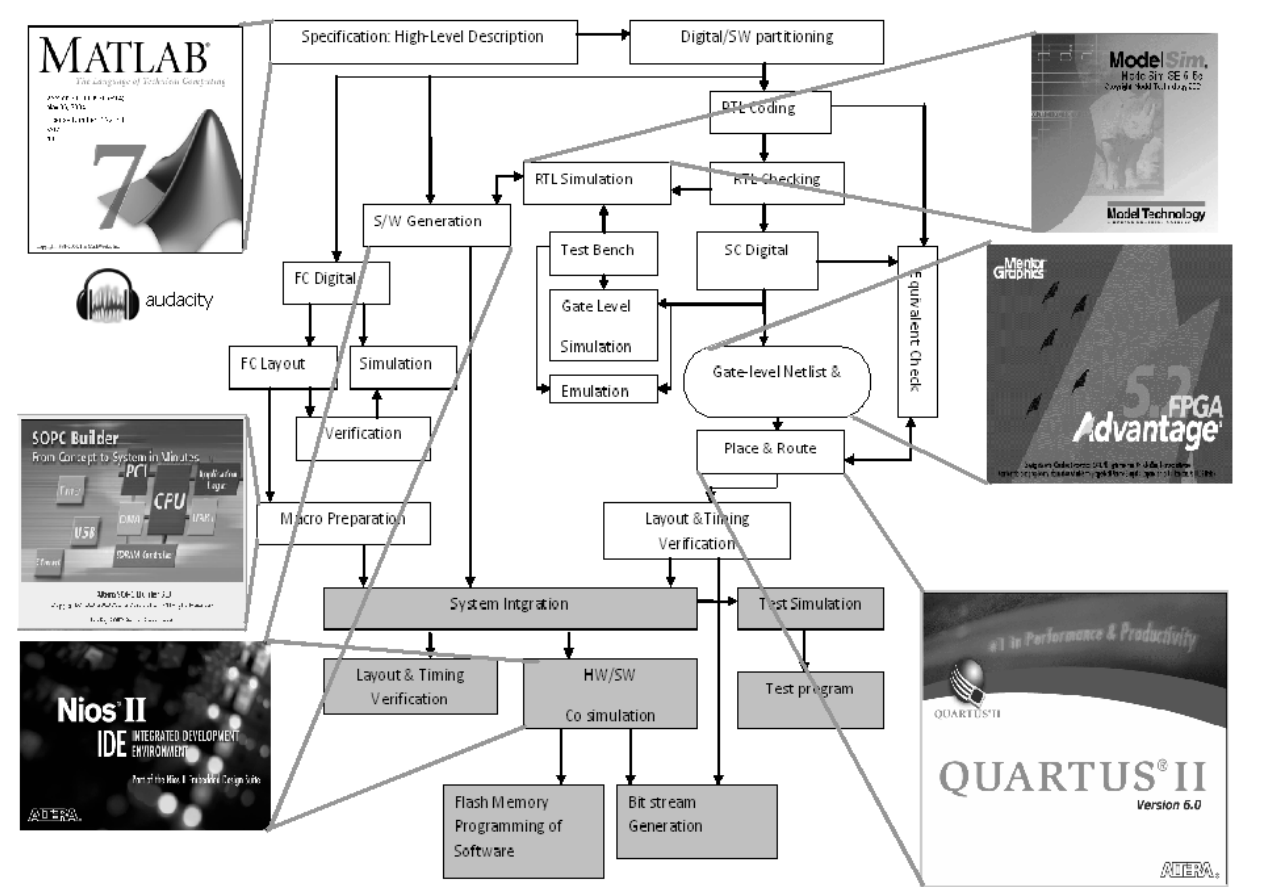


Fig. 13. Hardware/Software flow of the IWR system Development

The following table depicts the low hardware requirement for the Filter bank based feature extraction unit which is implemented as a hardware-software combined block compared to a fully hardware LPC block and hence this implementation has very low power consumption.

Hardware Requirements	Hardware LPC Implementation	Nios 2 Processor system Requirements	SOC based IWR with Nios processor(parallel)
Total logic elements	29799	4900	18393
Total registers	221	3069	3704
Embedded Multiplier 9-bit elements	52	4	52
Total memory bits	0	79360	90880

Table 3. Comparison of Hardware Logic requirements for2 different Front End processing unit

The hardware is interfaced to software unit as a custom instruction in Altera FPGA and the recognition Viterbi algorithm is implemented as a parallel hardware as shown in figure .The overall chip architecture for the speech recognition which does noise robust processing is shown in fig along with the ALTERA FPGA resource utilization values.

6. Application of matrix processor to speech recognition

6.1 Signal subspace based speech enhancement

From the noisy speech samples, H_y is constructed from a frame of noisy speech vector $[y(0),y(1)...y(L-1)]^T$.The relationship between noisy speech and clean speech vectors can be denoted as $H_y= H_x+ H_n$ where H_n represents noisy speech vectors which are estimated during voice activity detection process [1]. SVD of a real matrix can be divided in matrix form as,

$$Hy = [U_{y1} \ U_{y2}] \begin{pmatrix} \Sigma_{y1} & 0 \\ 0 & \Sigma_{y2} \end{pmatrix} \begin{bmatrix} V_{y1}^T \\ V_{y2}^T \end{bmatrix}$$

(24)

If the noise vectors are considered white then the smallest eigen value of Σ_{y1} is significantly greater than largest eigen value of Σ_{y2} . One can use of the SVD reduction techniques such as Thin SVD,Truncated SVD ,Thick SVD for signal space reduction and this implementation considers major p eigen values for signal subspace reduction based speech enhancement[5].

Frobenius norm constraint

Reconstructed signal is represented as a linear combination of the eigen vectors corresponding to the most significant eigen values and the dimensions of the hankel matrix is considered based on the Frobenius norm based performance criteria.

$$\overline{H_x} = U_{y1} \Sigma_{y1}^{-1} (\Sigma_{y1}^2 - \sigma_w^2 I) V_{y1}^T$$

(25)

$$\phi^{(s)} = \left\| H_y \right\|_F^2 - \left\| H_x^{(s)} \right\|_F^2 - \left\| H_w \right\|_F^2$$

(26)

Where $\phi^{(s)}$ represents the error associated with speech enhancement considering most p significant Eigen values and the suffix F indicates the Frobenius norm constraint of the

Hankel matrix $\|H_x^{(s)}\|_F^2$ represents the frobenius norm of the reconstructed clean speech matrix with p dominant eigen values and $\|H_w\|_F^2$ represents the Frobenius norm of the noise which is computed during Voice Activity Detection stage. Though progressive implementation of SVD helps to find the optimum p eigen values, real time implementation of such a system is complex intems of computational complexity and latency. Hence the approximate number of dominant eigen values is found to be 12 for speech samples present in AN4 database using MATLAB and is used for this speech enhancement application.

7. The matrix processor

This matrix processor implements three basic Matrix operations namely Addition/Subtraction, Fast Matrix Multiplication and fast Matrix Transposition using efficient Address generation unit in $O(n \log_2 n)$ operations and complex Matrix operations namely Singular Value Decomposition, QR decomposition, Matrix Bi diagonalization and it also implements matrix inversion for toeplitz matrices through Levinson-Durbin Algorithm.

7.1 Matrix multiplication

Matrix multiplication is done using Strassen's algorithm for matrix multiplication which is a superior algorithm for matrix multiplication of higher order as the computational complexity is $O(n^{2.81})$ compared to the normal row by row multiplication which requires $O(n^3)$ [14]. Any matrix A can be sub divided into block matrices and the product $C=A*B$ can be implemented using Strassen's algorithm at block level.

$$A = \begin{pmatrix} A_{11} & A_{12} \\ A_{21} & A_{22} \end{pmatrix} \quad (27)$$

The number of multiplications involved in strassen algorithm of block matrix multiplication is $(\frac{7}{8}n + \frac{3}{2})n^2$ and in strassen's algorithm the number of multiplications decreases as n increases whereas the number of addition remains same[5].

7.2 Concept of orthogonal matrix:

If vector product of columns of a matrix yields then the matrix is called as Orthogonal matrix. Thus for an orthogonal matrix inverse can be computed by simply taking the transpose of the matrix. The eigwn values of a matrix can be classified as r nonzero eigenvalues and $n - r$ zero eigenvalues. It is common convention to order the eigenvalues so that

$$\lambda_{1(\max)} \geq \lambda_2 \geq \dots \geq \lambda_{r(\min \text{ nonzero})} \gg \lambda_{r+1} > \lambda_n$$

that λ_1 is the largest, with the remaining nonzero eigenvalues arranged in descending order, followed by $n - r$ zero eigenvalues. Note that if A is full rank, then $r = n$ and there are no zero eigenvalues. The quantity λ_n is the eigenvalue with the lowest value.

7.2 Givens rotation

Our Hardware implementation of QR decomposition and Singular value decomposition are based on Gentleman kung parallel systolic architecture which uses $O(m+n)$ complexity as

$A \in R^{m \times n}$ matrices where as the other implementations require $3n^2(m - \frac{n}{3})$ FLOPS to

compute QR decomposition (Press 1992). Two kinds of Processing Elements were used in the design of SVD, QR decomposition and matrix bi-diagonalization blocks namely Vectoring Mode CORDIC Processing Element and Rotation Mode CORDIC Processing Element. Since any orthogonal matrix can be resolved into several rotation matrices using Euler-Brauer Resolutions, SVD factorization using givens rotations suits hardware implementation (refer Figure 14).

7.3 The QR-decomposition

QR decomposition of a matrix yields an unitary matrix Q whose columns are orthogonal to each other and an upper triangular matrix R . Some characteristics like the l_2 -norm of a vector remain unaltered after QR decomposition. Householder reduction is one numerically stable method for the computation of QR decomposition. If a vector lies in the orthogonal complement subspace then its range of vector space is zero and hence it can be used to nullify some of the elements in a matrix. In the linear equation, elimination of one variable leads to a reduced space which can be done by choosing a vector so that the reflection of it in span of a particular vector space lines up with the x-axis. Q_{pq} is an orthonormal matrix which on recursively operated on the data matrix A , results in triangular matrix $A1$ as $A_1 = Q_{12}Q_{13}Q_{14} \dots A^T$. This processing matrix is post multiplied on the data matrix A using the Givens rotation algorithm mentioned as above. One rotation premultiplication by Q_{pq} exists for every element to be eliminated. $\square$

$$Q_{pq} = \begin{pmatrix} 1 & & \dots & & 0 \\ & \ddots & & & \\ & & \begin{pmatrix} \cos_p \theta & \sin_q \theta \\ -\sin_p \theta & \cos_q \theta \end{pmatrix} & & \\ & & & \ddots & \\ 0 & & & & 1 \end{pmatrix} \quad (28)$$

The following conditions needs to be met to eliminate the row elements :

1. If a matrix A is to be orthonormal, then each J must be orthonormal. (Because the product of orthonormal matrices is orthomonal).
2. The (I,k) th element JA must be zero.

First condition is statisfied by the following relation given below.

$$Q^T Q = \begin{pmatrix} 1 & & \dots & & 0 \\ & \ddots & & & \\ & & \begin{pmatrix} \cos_p \theta & -\sin_q \theta \\ \sin_p \theta & \cos_q \theta \end{pmatrix} & & \\ & & & \ddots & \\ 0 & & & & 1 \end{pmatrix} \begin{pmatrix} 1 & & \dots & & 0 \\ & \ddots & & & \\ & & \begin{pmatrix} \cos_p \theta & \sin_q \theta \\ -\sin_p \theta & \cos_q \theta \end{pmatrix} & & \\ & & & \ddots & \\ 0 & & & & 1 \end{pmatrix} = \begin{pmatrix} 1 & & \dots & & 0 \\ & \ddots & & & \\ & & \begin{pmatrix} 1 & \\ 0 & 1 \end{pmatrix} & & \\ & & & \ddots & \\ 0 & & & & 1 \end{pmatrix} \quad (29)$$

The second constraint is given by ,

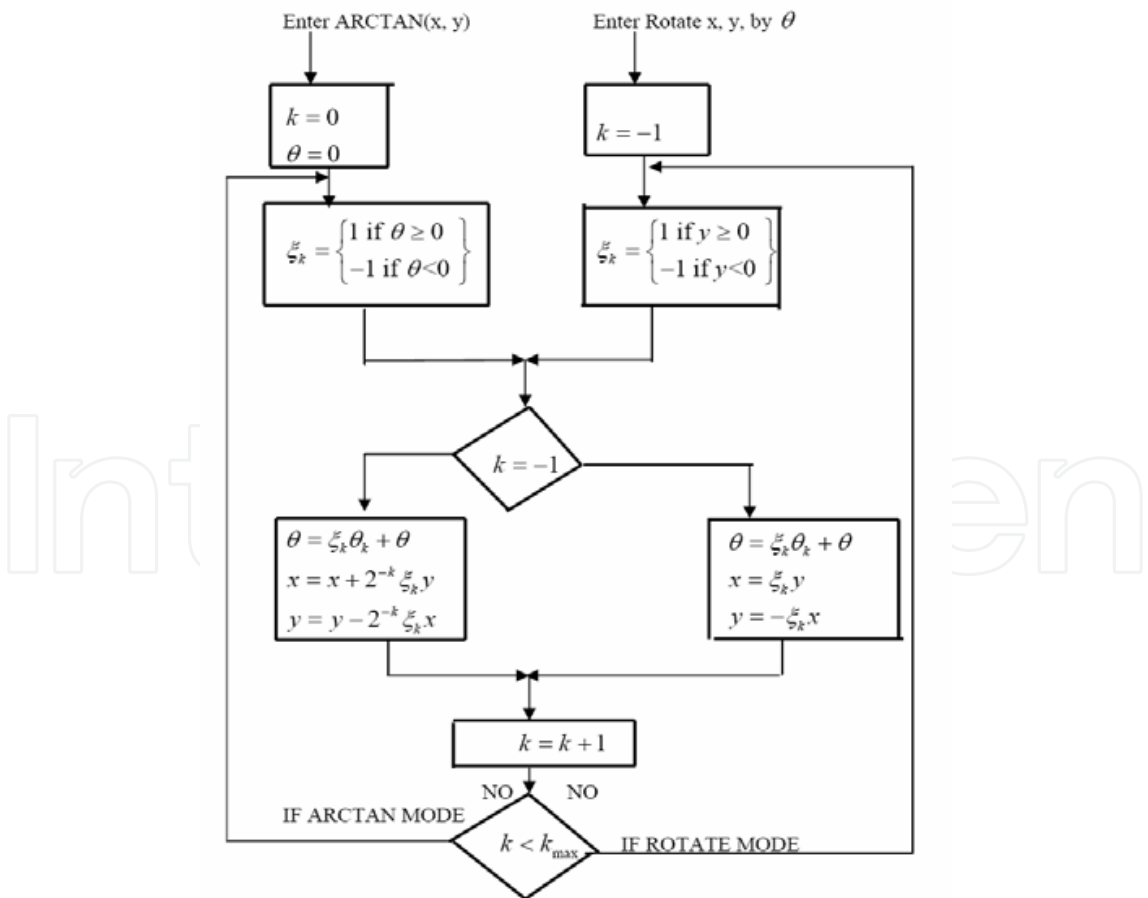
$$QA = \begin{pmatrix} 1 & \dots & 0 \\ \vdots & \ddots & \vdots \\ 0 & \dots & 1 \end{pmatrix} \begin{pmatrix} a_{11} & \dots & a_{1n} \\ \vdots & \ddots & \vdots \\ a_{n1} & \dots & a_{nn} \end{pmatrix} \quad (30)$$

which evaluates the following equations as,

$$\begin{aligned} -\sin_p \theta * a_{kk} + \cos_q \theta * a_{ik} &= 0 \\ \tan \theta &= \frac{a_{ik}}{a_{kk}} \end{aligned} \quad (31)$$

$$\sin_p \theta = \frac{a_{ik}}{\sqrt{a_{ik}^2 + a_{kk}^2}} \quad (32)$$

$$\cos \theta = \frac{a_{kk}}{\sqrt{a_{ik}^2 + a_{kk}^2}}$$



A fast algorithm for evaluating the above equations is given as follows.
Fig. 14. Flow chart for Givens Rotation algorithm

7.4 The singular value decomposition

The SVD is a method of extracting a diagonal matrix from a real matrix A which satisfies the following relationship $A = UDV^T$ where U and V are unitary matrices. SVD of any matrix is computed by decomposing the given matrices to 2×2 smaller dimension matrices and taking its inverse which is described in (Press 1992). This implementation considers a two step two sided unitary transformation in which each step is a diagonalization step (Hemkumar 1991) on in which each step is a diagonalization.

7.5 Systolic array based singular value decomposition for inverse computation:

The basis of the systolic array is the processing element PE. The PE is a simple computational devise capable of performing basic multiply and accumulate operation. At the beginning of each clock cycle,the PE reads in the values A_{ij} and C_k performs the necessary arithmetic computation using the internally stored value. By this above method only one PE is busy at a time and only one row of matrix can be considered at a time. It is possible to evaluate all elements of the matrix using a vector C concurrently with processors being busy all the time. All that is necessary to place n rows of processing elements beneath the first row. Then the computation of the second inner product involving the second row of A follows directly behind the computation of the first element of the product, similarly with third row etc. After m clock periods the m th row begins to accumulate and after m cycles the results will become stable and can be stored in the output register array. The basic advantage of using a systolic array in the matrix computation is that these computational blocks are regular. The PE's only talk to nearest neighbors. The above points make the silicon VLSI layout of this computational structure relatively simple. Only one cell need be designed the entire array is formed by repeating this design many times which is a simple process in VLSI design. The interconnections between processors are simple because they talk only to nearest neighbors.The idea behind this systolic array is to produce a massively parallel computational architecture which is capable of executing the QR decomposition in $O(m \cdot n)$ time units. As we see later conventional implementations require $O(3n^2(m-n))$ units

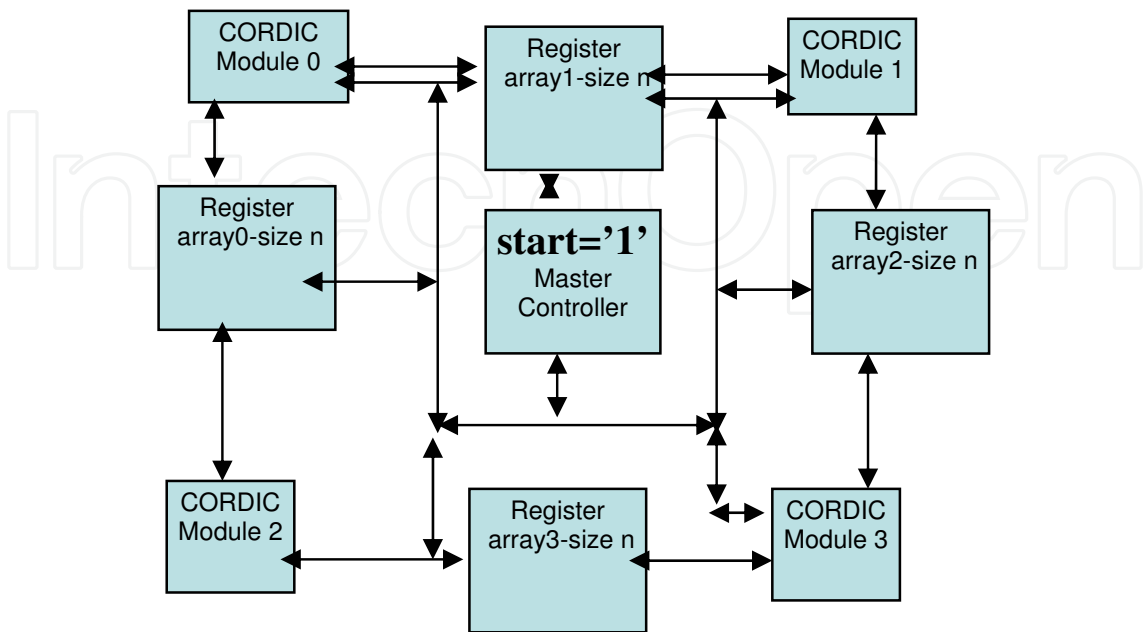


Fig. 15. SYSTOLIC CORDIC ARCHITECTURE FOR SVD

to compute the QR decomposition using the Givens procedure. Thus systolic architecture can be much faster. Further this procedure avoids the inefficiency of having to calculate the entire solution over again from the start for each iteration. Thus this systolic array gives us a new form of adaptive iterative structure, which can track changes in the LS solution as the environment changes or evolves. We now consider the systolic structure to eliminate the element a .

Two sided Jacobi algorithm is used for the computation of eigenvalue / singular value decomposition of a general matrix (Figure 15) A. As the serial jacobi algorithm is slow, a parallel systolic architecture is used with dynamic parallel ordering. This approach reduces the number of outer parallel iterations steps that two sided jacobi algorithm by 30-40 percent and hence It speeds up the inverse computation step. Such pre-conditioning should concentrate the Frobenius norm of A to diagonal as much as possible. For a diagonal matrix only one outer parallel iteration step would be required for the whole SVD computation and hence the concentration of Frobenius norm towards the diagonal might decrease the number of outer parallel iterations substantially.

7.5 Levinson-Durbin toeplitz solver

If the input matrix is symmetric as well as toeplitz system then by applying Levinson and Durbin's Algorithm one can solve the equations. Levinson's algorithm is an iterative algorithm and can be used to get only partial result for a toeplitz matrix and if the input matrix is also symmetric then one can make use of the Durbin's recursion to carry through the partial result to the final solution. Levinson-Durbin algorithm represents recursive system which solves toeplitz matrices and can be considered to model Auto Regressive models progressively and based on some criteria like AIC model order selection criteria one can choose the optimal model using LD algorithm. Normal matrix inverse computation requires $O(n^3)$ computations which can be done through Gaussian Elimination. If the matrix has got special structure as that of Toeplitz the inverse of the system can be found by this faster algorithm which models the linear set of equations to that of an auto regressive filter. Set of equations that describe Levinson- Durbin algorithm is given below. At the r th step of the LD Algorithm solves the r th truncated problem of Auto regressive filter design.

$$\begin{pmatrix} r(0) & \dots & r(M) \\ \vdots & \ddots & \vdots \\ r(-M) & \dots & r(0) \end{pmatrix} \begin{pmatrix} a_{M,0} \\ \vdots \\ a_{M,M} \end{pmatrix} = \begin{pmatrix} \rho_M \\ 0 \\ \vdots \end{pmatrix} \quad (33)$$

7.6 Fast matrix transposition unit

Large data matrices can be process efficiently can be loaded into memory using parallel architecture and for a n by n array stored in memory, it is convenient to transfer one column at a time and transposing a matrix means swapping elements of two different columns which is indeed difficult if one uses the transposition by column by column access method. Transposition can be done by assuming block matrices as shown below.

$$A = \begin{pmatrix} A_{11} & A_{12} \\ A_{21} & A_{22} \end{pmatrix} \quad A^T = \begin{pmatrix} A_{11}^T & A_{21}^T \\ A_{12}^T & A_{22}^T \end{pmatrix} \quad (34)$$

By transposing only Totally n^2 columns must be read which necessitates $O(n^2)$ operations which can be reduces to $O(n\log_2n)$ by interchanging $A^{T_{21}}, A^{T_{12}}$ block transpose matrices. Recursively computing the four block metrics $A^{T_{11}}, A^{T_{12}}, A^{T_{21}}, A^{T_{22}}$ logarithmic complexity is obtained.

8. Hardware architecture

The system is implemented in ALTERA’S EP2C20f484C7 FPGA with NIOS 2 processor. Matlab tool was used for designing system level specification and the RTL simulation is carried out in Modelsim. So as to reduce energy consumption over the recognition period, the operating frequency of the system is set to 12.5 MHz compared to the maximum operating frequency of 33 MHz, which results in low power consumption of 32mW which is 23% decrease in power compared to the previous case. For ASIC implementation, RTL compiler is used for synthesis and advanced Synthesis flows like Multi-Vt, DFT have been exercised in our design and the results are shown in Table 1.

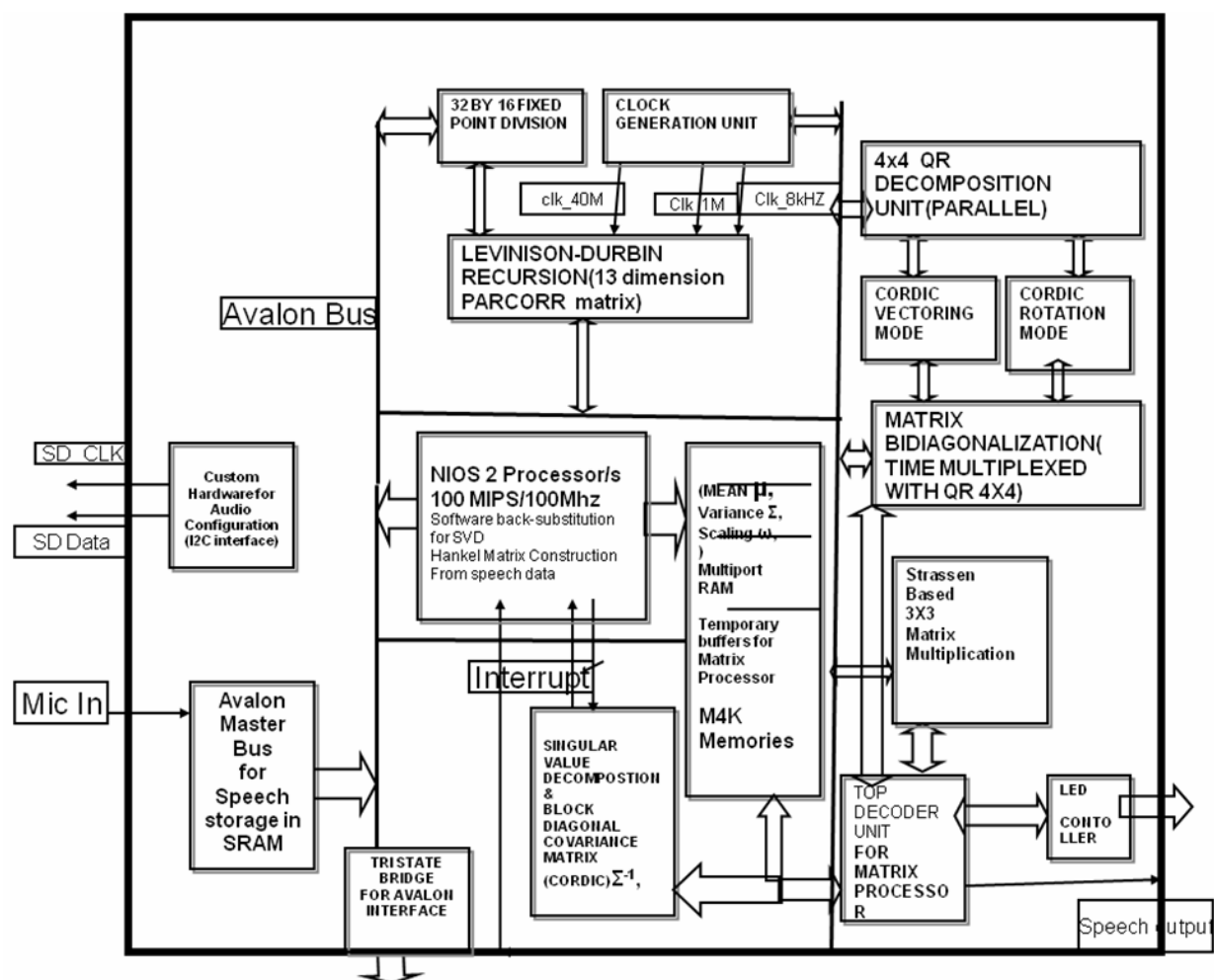


Fig. 16. Overall Hardware Architecture in Altera FPGA.

The hardware module is interfaced to software unit as a custom instruction in Altera FPGA and the latency of individual units are shown in Table2. AN4 database from Carnegie Mellon University was used to train the word models. Hardware emulation of the recognizer indicates a 5% improvement in speech recognition in noisy environment at the cost of 10%

increase in hardware complexity. Since the SVD block is used in time multiplexed fashion, 15% improvement in recognition accuracy can be obtained with the proposed unit. The following table depicts the low hardware requirement for the Filter bank based feature extraction unit which is implemented as a hardware-software combined block compared to a fully hardware LPC block and hence this implementation has very low power consumption. The following tables shows the ASIC implementation results for the matrix processor block and its power consumption values. This particular block can operate at a maximum frequency of 1Ghz. Complete power estimation for this chip was not carried out due to unavailability of the IO book power information and processor power dissipation information.

Module	Power(nW)	Area (micro m ²)	Gates	Frequency MHz
SVD	2431120	25320	3685	1000
QR Decomposition	2005727	27109	5398	1000
Levinison-Durbin	34276231	106519	8032	250
CORDIC Rotation	571500	5308	758	1000
CORDIC Vectoring	241919	4557	632	1000

Table 4. ASIC implementation results for matrix processor

9. Speech recognition results

No. of States		K=4	K=6	K=8	K=10	K=16	K=32
Memory Required for Feature Vectors - 13 Dimensions (in FPGA resources)	Covariance	910	910	910	910	910	910
	Mean	3900	5850	7800	9750	15600	31200
	Transition matri	1280	2880	5120	8000	20480	81920
	Initial probability	58	87	115	184	230	460
	Total Memory	6148	9727	13945	18844	37220	114490
Memory Required 26 Dimensions (in FPGA resources)	Covariance	1820	1820	1820	1820	1820	1820
	Mean	7800	11700	15600	19500	31200	62400
	Transition matrix	1280	2880	5120	8000	20480	81920
	Initial probability	58	87	115	184	230	460
	Total Memory	10958	16487	22655	29504	53730	145600
Memory Required 39 Dimensions (in FPGA resources)	Covariance	2730	2730	2730	2730	2730	2730
	Mean	11700	17550	23400	29250	46800	93600
	Transition matri	1280	2880	5120	8000	20480	81920
	Initial probability	58	87	115	184	230	460
	Total Memory	15768	23247	31365	40164	70240	178710

Table 5. Logic Utilization for realizing Memory in FPGA for various HMM states

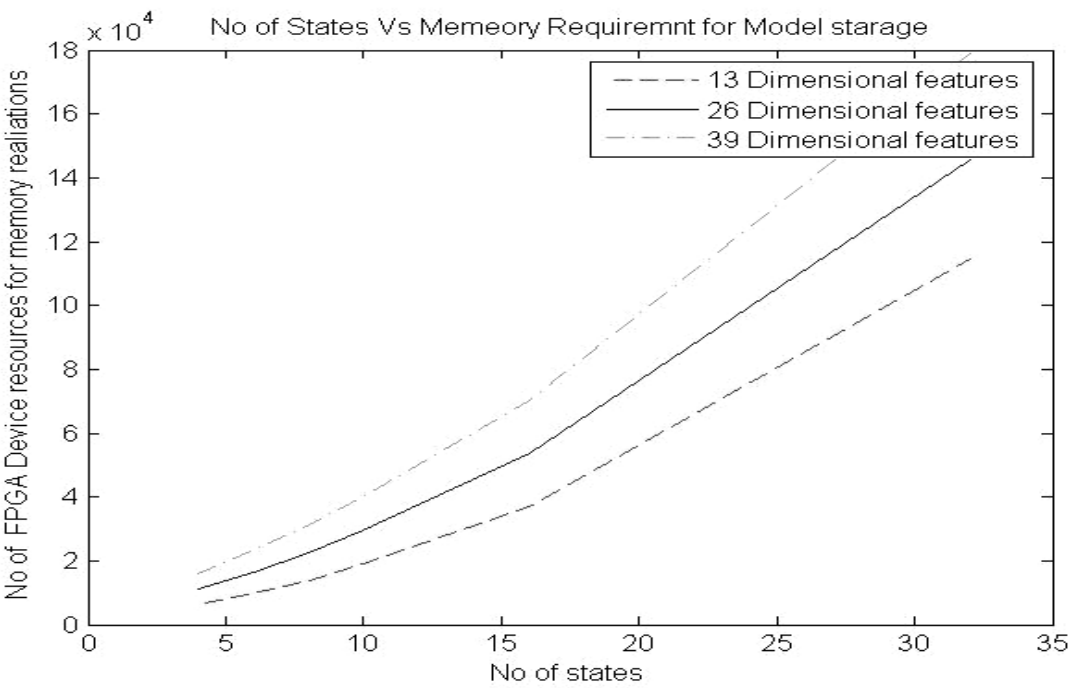


Fig. 16. Logic Utilization for realizing Memory in FPGA for various HMM states

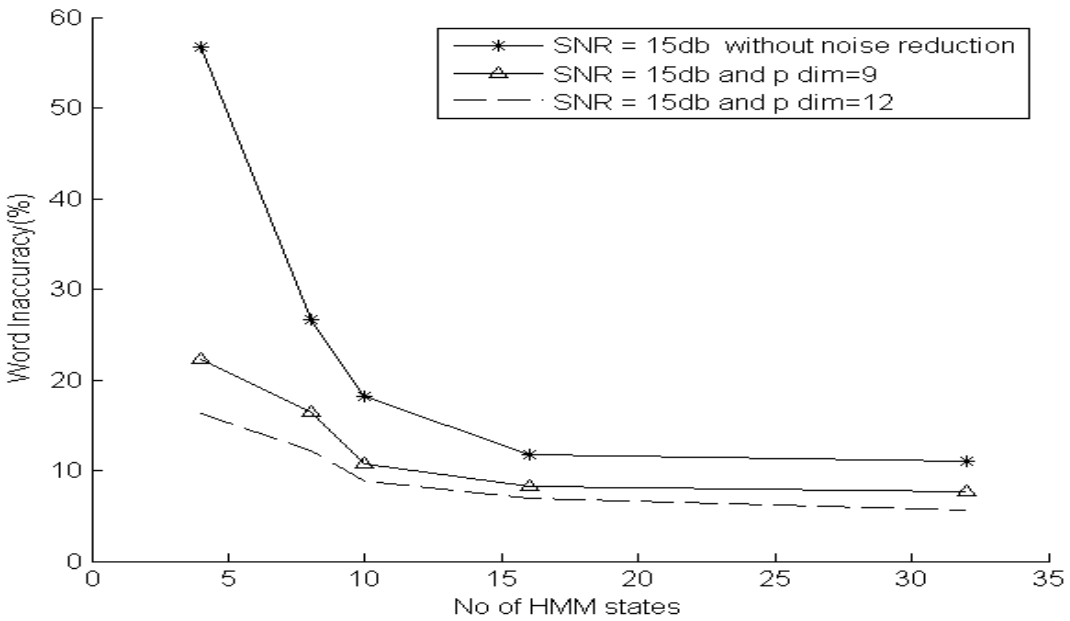


Fig. 17. Word Error rate for Enhanced speech for different SNR conditions.

AN4 database from Carnegie Mellon University has been used to train the word models. AN4 Database has got 130 different words spoken by people of different ages, dialect and gender. As the time complexity for the training of 130 words is extremely high, 10 digits from these 130 words have been used for training of the database and the system has been developed for 10 words. The AN4 database files are in raw PCM format, sampled at 16 kHz, in big endian byte order. Since the software IWR system is developed for 8 KHz sampling rate, A tool called Audacity is used to convert the 16Khz raw data file to 8Khz Wav file which was later fed to Matlab to extract the feature vectors using the feature extraction algorithm. The recognition performance of the proposed IWR system for different Number of states is plotted in Figure 16 & Figure 17.

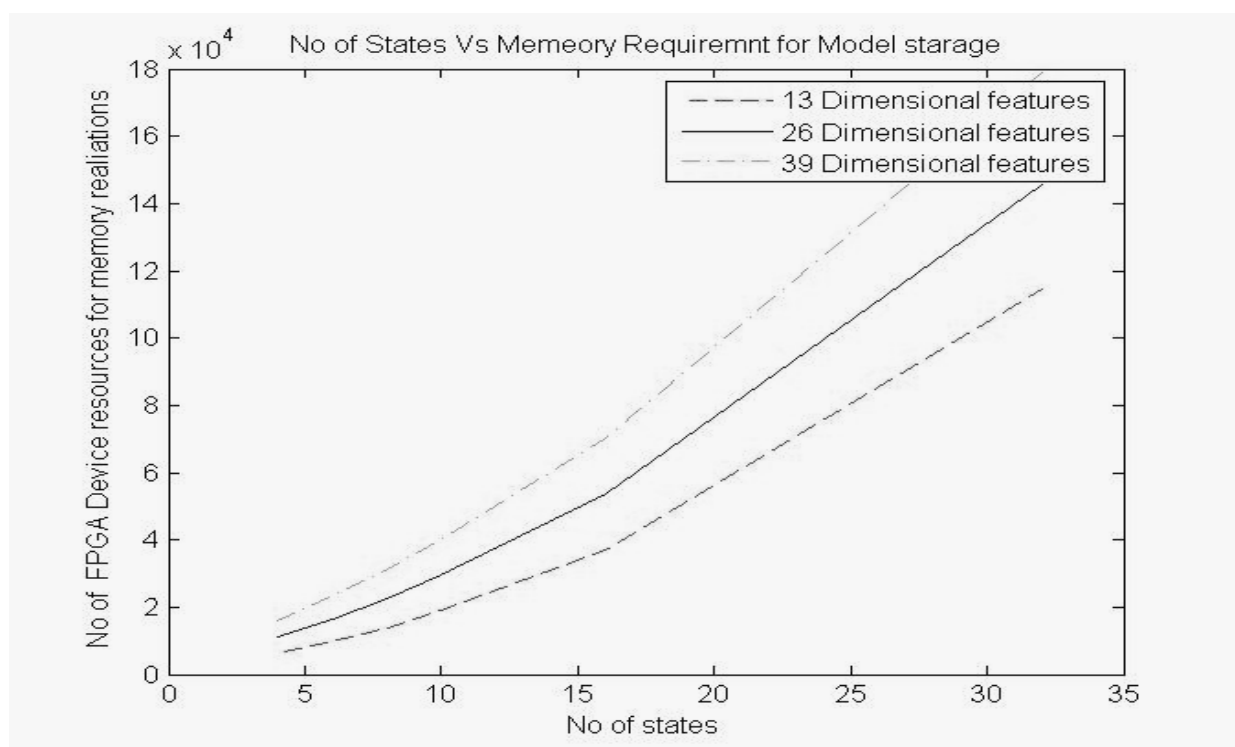


Fig. 18. Memory requirement for HMM models with different states and feature vectors

10. Conclusion

This 32 state Continuous Hidden Markov Model based speech recognition hardware provides 82.7% recognition accuracy in noisy speech environments at 15db for a 10 word sample space collected from words uttered by people of different age, gender and dialect as the data are processed in identical fashion. Hardware emulation of the recognizer indicates a 5% improvement in speech recognition in noisy environment at the cost of 10% increase in hardware complexity. Since the SVD block is used in time multiplexed fashion, 15% improvement in the overall recognition accuracy has been obtained.

ALU OPTION	Hardware Details	Cycles per Instruction	Result Latency Cycles	Parallel implementation of the 13 features with 4 states
LE Based Multiplier	32 x 4-bit multiplier	11	2	76
Embedded Multiplier on Cyclone	ALU includes 32 x 16-bit multiplier	5	2	24
Hardware Divide	ALU includes multicycle divide	4-66	2	0
Custom Avalon Speech Controller	Custom Instruction1	2	2	1
Custom Avalon SVD	Custom Instruction2	16	2	1

Table 6. Hardware Execution cycles of the Custom Instructions in NIOS Processor

11. References

H.Abut, H.L. Hansen, K.Takeda, "DSP for In Vehicle and Mobile Systems", Springer Publishers, 2005.

J. Pihl, T. Svendsen, and M. H. Johnsen, "A VLSI implementation of pdf computations in HMM based speech recognition," in Proc. IEEE TENCON'96, 1996, pp. 241-246.

K. Hermus, P.Wambacq, and H.V. Hamme, "A Review of Signal Subspace Speech Enhancement and Its Application to Noise Robust Speech Recognition", EURASIP Journal on Advances in Signal Processing Volume 2007, Article ID 45821.

L.R.Rabiner & B.H. Juang," Fundamentals Of Speech Recognition." Prentice -Hall, AT&T, U.S.A, 1993.

Simon Haykin, "Adaptive Filter Theory", Third Edition , Prentice Hall Information and System series,2002.

N. D. Hemkumar, "A Systolic VLSI Architecture for Complex SVD", Postgraduate Thesis, Rice University, Houston, Texas, May 1991.

N.Karthikeyan, S.Arun, K.Murugarraj, M.John, "An application specific matrix processor for signal subspace based speech enhancement in noise robust speech recognition applications", pages 766-769,7th Internation Conference on ASIC(ASICON2007), Guilin, China,2007.

N.Karthikeyan, S.Arun, K.Murugaraj, M John, "Hardware and Software acceleration of Front End Processing Unit in Scalable Noise robust Word HMM based speech recognition systems", 4th International Conference on SOC (ISOCC-2007), Seoul, Korea, (Poster).

S. Nedeveschi, R.K. Patra, E.A.Brewer,"Hardware Speech Recognition for User Interfaces in Low Cost, Low Power Devices", DAC 2005, June 13.17, 2005, Anaheim, California, USA.

- S.Yoshizawa, N. Wada, N. Hayasaka, IEEE, and Yoshikazu Miyanaga, "Scalable Architecture for Word HMM-Based Speech Recognition and VLSI Implementation in Complete System", *IEEE TRANSACTIONS ON CIRCUITS AND SYSTEMS – I: REGULAR PAPERS*, VOL. 53, NO. 1, JANUARY 2006
- Saeed V. Vaseghi, "Advanced Digital Processing and Noise Reduction", John Wiley and Sons Publishers, Second Edition, 2000.
- U. Mayer Baese, "Digital Signal Processing with Field Programmable Gate Arrays", Springer Publishers, 2005.
- William H. Press, Saul A. Teukolsky, William T. Vetterling and Brian P. Flannery "Numerical Recipes in C", CAMBRIDGE UNIVERSITY PRESS, 1992.
- "Avalon Streaming Interface Specification", Version 1.0, Altera Corporation, November 2006.

IntechOpen



Speech Recognition

Edited by France Mihelic and Janez Zibert

ISBN 978-953-7619-29-9

Hard cover, 550 pages

Publisher InTech

Published online 01, November, 2008

Published in print edition November, 2008

Chapters in the first part of the book cover all the essential speech processing techniques for building robust, automatic speech recognition systems: the representation for speech signals and the methods for speech-features extraction, acoustic and language modeling, efficient algorithms for searching the hypothesis space, and multimodal approaches to speech recognition. The last part of the book is devoted to other speech processing applications that can use the information from automatic speech recognition for speaker identification and tracking, for prosody modeling in emotion-detection systems and in other speech processing applications that are able to operate in real-world environments, like mobile communication services and smart homes.

How to reference

In order to correctly reference this scholarly work, feel free to copy and paste the following:

Karthikeyan Natarajan, Mala John and Arun Selvaraj (2008). Silicon Technologies for Speaker Independent Speech Processing and Recognition Systems in Noisy Environments, *Speech Recognition*, France Mihelic and Janez Zibert (Ed.), ISBN: 978-953-7619-29-9, InTech, Available from:

http://www.intechopen.com/books/speech_recognition/silicon_technologies_for_speaker_independent_speech_processing_and_recognition_systems_in_noisy_envi

INTECH
open science | open minds

InTech Europe

University Campus STeP Ri
Slavka Krautzeka 83/A
51000 Rijeka, Croatia
Phone: +385 (51) 770 447
Fax: +385 (51) 686 166
www.intechopen.com

InTech China

Unit 405, Office Block, Hotel Equatorial Shanghai
No.65, Yan An Road (West), Shanghai, 200040, China
中国上海市延安西路65号上海国际贵都大饭店办公楼405单元
Phone: +86-21-62489820
Fax: +86-21-62489821

© 2008 The Author(s). Licensee IntechOpen. This chapter is distributed under the terms of the [Creative Commons Attribution-NonCommercial-ShareAlike-3.0 License](https://creativecommons.org/licenses/by-nc-sa/3.0/), which permits use, distribution and reproduction for non-commercial purposes, provided the original is properly cited and derivative works building on this content are distributed under the same license.

IntechOpen

IntechOpen