



Literature Review on Speech Based Information Retrieval through web for visually impaired

Ananthi. S*

Ph.D Scholar, Dept. of CSE
Annamalai University
Chidambaram, India.
ananthi68@gmail.com

Dhanalakshmi. P

Associate Professor, Dept. of CSE
Annamalai University
Chidambaram, India.
abi_dhana@rediffmail.com

Abstract: Speech-based content retrieval system facilitates the visually impaired to have an efficient access, search and browsing capabilities to the desired content in audio form. In Speech based content retrieval method, Query is given as voice command to the system; the system has to generate the information based on the spoken word. The information has to be retrieved based on the keyword matching. The resultant information should be converted as a speech. Usual Method used in Speech Recognition (SR) is Hidden Markov Model (HMM), Dynamic Time Warping (DTW) [1]. If the generated system is Speaker Dependent, then it can't analyze the user's voice (input) for further process. So the system has to be implemented for Speaker Independent. So that it can generate speech signals and its corresponding text [2]. People with visual disability or impaired can make this type of web search.

Keywords: Dynamic Time Warping (DTW), Hidden Markov Model (HMM), Hyperlinks and Algorithm, Information Retrieval, Keyword Matching, Rough-fuzzy Method, Speech Recognition (SR), Speech Synthesis.

I. INTRODUCTION

Users needs and user's request is given as audio Query which is converted into text, based on Speech Recognition. The converted text is given as input to the system to retrieve the information based on keyword matching. This is achieved by Text to Text Conversion technique. Retrieved text is again converted into Speech, based on Speech Synthesis Technique. Speech Recognition is the process of analyzing the entire phrase in order to promptly and effectively provides the appropriate outcome [3]. Speech recognition (SR) is the translation of spoken words into text. It is also known as automatic speech recognition (ASR), or computer speech recognition.

Some SR systems use "training" where an individual speaker reads sections of text into the SR system. These systems analyze the person's specific voice and use it to fine tune the recognition of that person's speech, resulting in more accurate transcription. Systems that do not use training are called "Speaker Independent" systems. Systems that use training are called "Speaker Dependent" systems [4]. The term voice recognition refers to finding the identity of "who" is speaking, rather than what they are saying. Recognizing the speaker can simplify the task of translating speech in systems that have been trained on specific person's voices or it can be used to authenticate or verify the identity of a speaker as part of a security process[3].

Speech Recognition Engine has the ability to answer questions and effectively provide appropriate answers as if it was an actual live human. Other technologies use simple key word detection. Natural Language Speech Recognition actually analyzes the entire phrase in order to promptly and effectively provide the appropriate outcome [4], [5].

Natural Language Speech Recognition has the ability to adapt to the callers request and make real time adjustments by using built in logic. This technology allows you to treat customer's requests 24\7.

Speech synthesis is the artificial production of human speech. A computer system used for this purpose is called a speech synthesizer, and can be implemented in software or hardware [5], [6]. Text-to-speech (TTS) systems convert normal language text into speech; other systems render symbolic linguistic representations like phonetic transcriptions into speech. Synthesized speech can be created by concatenating pieces of recorded speech that are stored in a database [6].

Systems differ in the size of the stored speech units; a system that stores phones or diaphones provides the largest output range, but may lack clarity. For specific usage domains, the storage of entire words or sentences allows for high-quality output. Alternatively, a synthesizer can incorporate a model of the vocal tract and other human voice characteristics to create a completely "synthetic" voice output.

In General there are three usual methods in speech recognition: Dynamic Time Warping (DWT), Hidden Markov Model (HMM) and Artificial Neural Networks (ANNs) [1].

Information Retrieval (IR) is a branch of Computer Science dealing with storage, maintenance and information search within large amounts of data. Information is retrieved through the relevant information from the web using text-to-text based on the entered key i.e., based on keyword matching technique. This is achieved by techniques like Rough fuzzy algorithm and web structure mining by Exploring Hyperlinks and Algorithm [7], [8].

II. EXISTING SYSTEM

A. Text based information retrieval:

The information is retrieved based on the typed text in search engines. Input query is given as text format typed by the user. To retrieved relevant information from the web using text-to-text keyword matching technique (information

retrieval). The retrieved information is displayed in the screen.

B. Demerit:

- Person with visual disability can't work with this type of information retrieval through the net.
- Old age persons can't visual the system for longer time.

III. SPEECH RECOGNITION TECHNIQUES

A. Dynamic time warping:

Dynamic time warping (DTW) is a technique that finds the optimal alignment between two time series if one time series may be warped non-linearly by stretching or shrinking it along its time axis [9]. This warping between two time series can then be used to find corresponding regions between the two time series [9], [10]. In Speech Recognition Dynamic time warping is often used to determine if two waveforms represent the same spoken phrase. This method is used for time adjustment of two words and estimation their adjustment of two words and estimation their difference [10]. In a speech waveform, the duration of each spoken sound and the interval between sounds are permitted to vary, but the overall speech must be similar.

B. Demerits of DTW:

Main Problem of this system is it can generate only little amount of learning words. The calculating rate of the signal is high and it requires large memory for storing the speech signal also to generate respective text [9].

IV. PROPOSED SYSTEM

This Dynamic Time Warping is replaced by Hidden Markov Model (HMM). This method requires only less amount of memory and the memory requirement is less.

A. Hidden Markov model:

Hidden Markov Models are finite automates, having a given number of states; passing from one state to another is made instantaneously at equally spaced time moment. At every pass from one state to another, the system generates observations, two processes are taking place: the transparent one, represented by the observations string (feature sequence), and the hidden one, which cannot be observed, represented by the state string [11], [12]. Main point of this method is timing sequence and comparing method [1].

A hidden Markov model (HMM) is a statistical Markov model in which the system being modeled is assumed to be a Markov process with unobserved (hidden) states. An HMM can be considered as the simplest dynamic Bayesian network. The mathematics behind the HMM was developed by L. E. Baum and coworkers. It is closely related to an earlier work on optimal nonlinear filtering problem (stochastic processes) by Ruslan L. Stratonovich, who was the first to describe the forward-backward procedure [11].

In a regular Markov model, the state is directly visible to the observer, and therefore the state transition probabilities are the only parameters. In a hidden Markov model, the state is not directly visible, but output, dependent on the state, is visible. Each state has a probability distribution over the possible output tokens. Therefore the sequence of tokens generated by an HMM gives some information about the

sequence of states. Note that the adjective 'hidden' refers to the state sequence through which the model passes, not to the parameters of the model; even if the model parameters are known exactly, the model is still 'hidden'[12],[13].

Hidden Markov models are especially known for their application in temporal pattern recognition such as speech, handwriting, gesture recognition, part-of-speech tagging, musical score following, partial discharges and bioinformatics[11],[12],[13]. A hidden Markov model can be considered a generalization of a mixture model where the hidden variables (or latent variables), which control the mixture component to be selected for each observation, are related through a Markov process rather than independent of each other [12].

a. Description in terms of urns:

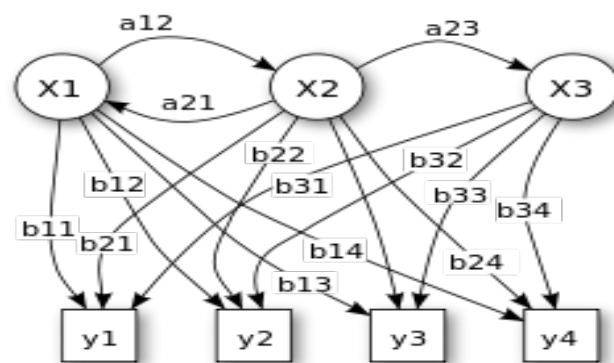


Figure 1. Probabilistic parameters of a hidden Markov model (example)

Where,

x — states

y — possible observations

a — state transition probabilities

b — output probabilities

In its discrete form, a hidden Markov process can be visualized as a generalization of the Urn problem: A genie is in a room that is not visible to an observer. In this hidden room there are urns X1, X2, X3...each of which contains a known mix of balls, each ball labeled y1, y2, y3. The genie chooses an urn in that room and randomly draws a ball from that urn. It then puts the ball onto a conveyor belt, where the observer can observe the sequence of the balls but not the sequence of urns from which they were drawn. The genie has some procedure to choose urns; the choice of the urn for the n-th ball depends only upon a random number and the choice of the urn for the (n - 1)th ball. The choice of urn does not directly depend on the urns chosen before this single previous urn; therefore, this is called a Markov process. It can be described by the upper part of Fig. 1.

The Markov process itself cannot be observed, and only the sequence of labeled balls can be observed, thus this arrangement is called a "hidden Markov process", where one can see that balls y1, y2, y3, y4 can be drawn at each state. Even if the observer knows the composition of the urns and has just observed a sequence of three balls, e.g. y1, y2 and y3 on the conveyor belt, the observer still cannot be sure which urn (i.e., at which state) the genie has drawn the third ball from. However, the observer can work out other details, such as the identity of the urn the genie is most likely to have drawn the third ball from[11],[13].

b. Architecture:

Initialization:

$$\alpha_1(i) = p_i b_i(o(1)), i = 1, \dots, N \text{ ---- (1)}$$

Recursion:

$$\alpha_{t+1}(i) = \left[\sum_{j=1}^N \alpha_t(j) a_{ji} \right] b_i(o(t+1)) \text{ ---- (2)}$$

Where

$$i = 1, \dots, N, t = 1, \dots, T - 1$$

Fig. 2 shows the general architecture of an instantiated HMM. Each oval shape represents a random variable that can adopt any of a number of values. The random variable $x(t)$ is the hidden state at time t (with the model from the above diagram, $x(t) \in \{x_1, x_2, x_3\}$). The random variable $y(t)$ is the observation at time t (with $y(t) \in \{y_1, y_2, y_3, y_4\}$). The arrows in the diagram (often called a trellis diagram) denote conditional dependencies.

From Fig. 2, it is clear that the conditional probability distribution of the hidden variable $x(t)$ at time t , given the values of the hidden variable x at all times, depends only on the value of the hidden variable $x(t-1)$: the values at time $t-2$ and before have no influence. This is called the Markov property. Similarly, the value of the observed variable $y(t)$ only depends on the value of the hidden variable $x(t)$ (both at time t).

In the standard type of hidden Markov model considered here, the state space of the hidden variables is discrete, while the observations themselves can either be discrete (typically generated from a categorical distribution) or continuous (typically from a Gaussian distribution). The parameters of a hidden Markov model are of two types, transition probabilities and emission probabilities (also known as output probabilities). The transition probabilities control the way the hidden state at time t is chosen given the hidden state at time $t-1$.

The hidden state space is assumed to consist of one of N possible values, modeled as a categorical distribution. (See the section below on extensions for other possibilities.) This means that for each of the N possible states that a hidden variable at time t can be in, there is a transition probability from this state to each of the N possible states of the hidden variable at time $t+1$, for a total of N^2 transition probabilities. Note that the set of transition probabilities for transitions from any given state must sum to 1. Thus, the N^2 matrix of transition probabilities is a Markov matrix. Because any one transition probability can be determined once the others are known, there are a total of $N(N-1)$ transition parameters.

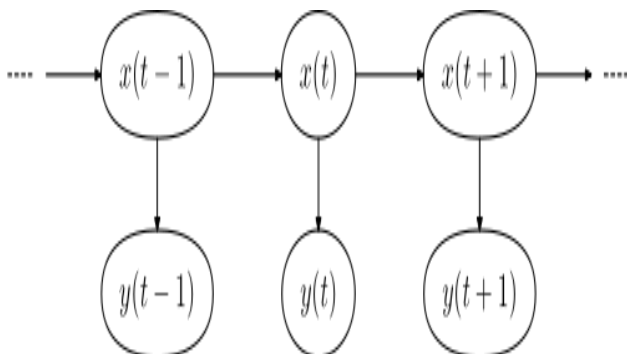


Figure 2. Transition Parameters

In addition, for each of the N possible states, there is a set of emission probabilities governing the distribution of the observed variable at a particular time given the state of the hidden variable at that time. The size of this set depends on the nature of the observed variable. For example, if the observed variable is discrete with M possible values, governed by a categorical distribution, there will be $M-1$ separate parameters, for a total of $N(M-1)$ emission parameters over all hidden states. On the other hand, if the observed variable is an M -dimensional vector distributed according to an arbitrary multivariate Gaussian distribution, there will be M parameters controlling the means and $M(M+1)/2$ parameters controlling the covariance matrix, for a total of emission parameters.

$$N(M + \frac{M(M+1)}{2}) = NM(M+3)/2 = O(NM^2) \text{ ---- (3)}$$

In such a case, unless the value of M is small, it may be more practical to restrict the nature of the co-variances between individual elements of the observation vector, e.g. by assuming that the elements are independent of each other, or less restrictively, are independent of all but a fixed number of adjacent elements).

B. Forward algorithm:

Let $\alpha_t(i)$ be the probability of the partial observation sequence $O_t = \{o(1), o(2), \dots, o(t)\}$ to be produced by all possible state sequences that end at the i^{th} state.

$$\alpha_t(i) = P(o(1), o(2), \dots, o(t) | q(t) = q_i) \text{ ---- (4)}$$

Then the unconditional probability of the partial observation sequence is the sum of $\alpha_t(i)$ over all N states. The Forward Algorithm is a recursive algorithm for calculating $\alpha_t(i)$ for the observation sequence of increasing length t . First, the probabilities for the single-symbol sequence are calculated as a product of initial i^{th} state probability and emission probability of the given symbol $o(1)$ in the i^{th} state.

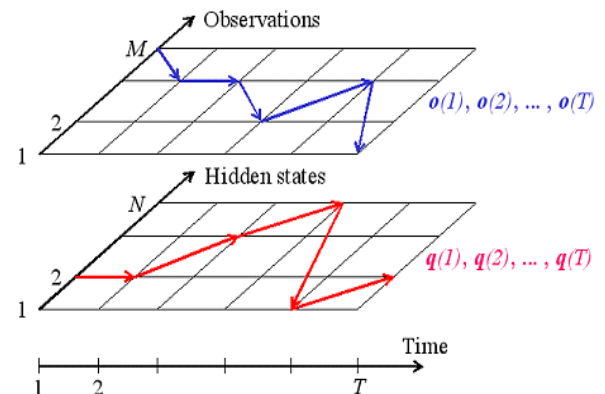


Figure 3. Observed and hidden sequences

Then the recursive formula is applied. Assume we have calculated $\alpha_t(i)$ for some t . To calculate $\alpha_{t+1}(j)$, we multiply every $\alpha_t(i)$ by the corresponding transition probability from the i^{th} state to the j^{th} state, sum the products over all states, and then multiply the result by the emission probability of the symbol $o(t+1)$. Iterating the process, we can eventually calculate $\alpha_T(i)$, and then summing them over all states, we can obtain the required probability.

Initialization:

$$\alpha_i(i) = p_i b_i(o(1)), i = 1, \dots, N \quad \text{----- (5)}$$

Recursion:

$$\alpha_{t+1}(i) = \left[\sum_{j=1}^N \alpha_t(j) a_{ji} \right] b_i(o(t+1)) \quad \text{----- (6)}$$

Where

$$i = 1, \dots, N, t = 1, \dots, T - 1 \quad \text{----- (7)}$$

Termination:

$$P(o(1)o(2) \dots o(T)) = \sum_{j=1}^N \alpha_T(j) \quad \text{----- (8)}$$

Algorithm:

Given A+E+(x1..., xn)

forward (x_seq, Q, fi, A, E):

T = {}

for q in Q:

T[q] = fi[q]

for x in x_seq:

U = {}

for q_next in X:

total = 0

for q in X:

total += T[q] * A[q][q_next] * E[q_next][x]

U[q_next] = total

T = U

V. INFORMATION RETRIEVAL TECHNIQUES

A. Rough fuzzy Method:

With the rapid growth of information on the World Wide Web, users normally face with large results returned from search engines with low precision. This is because the traditional search engine does not consider personal preferences in their searching algorithms and then information found by search engines is usually not relevant to a user's actual needs. Thus, personalized Web retrieval becomes increasingly important, which can conduct retrieval incorporating users' interested needs and help the user search information effectively and efficiently.

To achieve more effective retrieval for the user, a novel personalized Web retrieval approach based on rough-fuzzy hybridization is proposed. Variable precision rough set model (VPRSM) is used to deal with ambiguities of language and discovery user preference. To obtain relevance information, the user is asked to rank a set of training documents retrieved through a standard search engine. The user's preference is analyzed and discerning keywords are then identified to refine a query which represents the user's interests appropriately. In order to implement this process, fuzzy sets are incorporated to take care of the real-valued weights in each document vector.

Then, the refined query is fed to the search engine to retrieve better documents. To reflect the user's preferences further, retrieved results are re-ranked according to rough similarity measures between the refined query and documents.

B. Exploring Hyperlinks and Algorithm:

The Web is a massive, explosive, diverse, dynamic and mostly unstructured data repository, which delivers an incredible amount of information and also increases the complexity of dealing with the information from the different perspectives of knowledge seekers, Web service providers and business analysts. The following are considered as challenges (Da Gomes and Gong, 2005) in the Web mining:

- Web is huge and Web Pages are semi-structured.
- Web information tends to be diversity in meaning.
- Degree of quality of the information extracted
- Conclusion of the knowledge from the information extracted.
- Web mining techniques along with other areas like Database (DB), Information Retrieval (IR), Natural Language Processing (NLP) and machine learning can be used to solve the above challenges.

Web mining is the use of data mining techniques to automatically discover and extract information from the World Wide Web (WWW). Web structure mining helps the users to retrieve the relevant documents by analyzing the link structure of the Web. The basic problems in the Web structure mining and the mining categories are discussed below

a. Web structure mining

According to Kosala and Blockeel (2000), Web mining consists of the following tasks:

- Resource finding: The task of retrieving intended Web documents.
- Information selection and pre-processing: Automatically selecting and pre-processing specific information from retrieved Web resources
- Generalization: Automatically discovers general patterns at individual Web sites as well as across multiple sites.
- Analysis: Validation and interpretation of the mined patterns. There are three areas of Web mining according to the usage of the Web data used as input in the data mining process, namely, Web Content Mining (WCM).

Web content mining is concerned with the Retrieval of information from WWW into more structured forms and indexing the information to retrieve it quickly. Web usage mining is the process of identifying the browsing patterns by analyzing the user's navigational behavior. Web structure mining is to discover the model underlying the link structures of the Web pages, catalog them and generate information such as the similarity and relationship between them, taking advantage of their hyperlink topology. Hyperlink analysis and the algorithms discussed here are related to Web Structure mining. Even though there are three areas of Web mining, the differences between them are narrowing because they are all interconnected.

b. How big is web:

A Google report says that there are 1trillion (1,000,000,000,000) unique Universal Resource Locator (URLs) on the Web. The actual number could be more than that and Google could not index all the pages. When Google first created the index in 1998 there were 26 million pages and in 2000 Google index reached 1 billion pages. In the last 9 years, Web has grown tremendously and the usage of the

web is unimaginable. So it is important to understand and analyze the underlying data structure of the Web for effective Information Retrieval.

c. Web data structure:

The traditional information retrieval system basically focuses on information provided by the text of Web documents. Web mining technique provides additional information through hyperlinks where different documents are connected. The Web may be viewed as a directed labeled graph whose nodes are the documents or pages and the edges are the hyperlinks between them. This directed graph structure in the Web is called as Web Graph. A graph G consists of two sets V and E , Horowitz et al. (2008). The set V is a finite, nonempty set of vertices. The set E is a set of pairs of vertices; these pairs are called edges. The notation $V(G)$ and $E(G)$ represent the sets of vertices and edges, respectively of graph G . It can also be expressed $G = (V, E)$ to represent a graph. The considered graph is a directed graph with 3 Vertices and 3 edges.

A directed Graph (G) , vertices V of G $V(G) = \{A, B, C\}$.

The Edges E of G , $E(G) = \{(A, B), (B, A), (B, C)\}$. In a directed graph with n vertices, the maximum number of edges is $n(n-1)$. With 3 vertices, the maximum number of edges can be $3(3-1) = 6$. In the above example, there is no link from (C, B) , (A, C) and (C, A) . A directed graph is said to be strongly connected if for every pair of distinct vertices u and v in $V(G)$, there is a directed path from u to v and also from v to u . The graph in Fig. 3 is not strongly connected, as there is no path from vertex C to B . According to Broder et al. (2000), a Web can be imagined as a large graph containing several hundred million or billion of nodes or vertices and a few billion arcs or edges. The following paragraph explains the hyperlink analysis and the algorithms used in the hyperlink analysis for information retrieval.

d. Hyperlink analysis:

Many web pages do not include words that are descriptive of their basic purpose (for example rarely a search engine portal includes the word “search” in its home page) and there exist Web pages which contain very little text (such as image, music, video resources), making a text-based search techniques difficult. However, how others exemplify this page may be useful. This type of “characterization” is included in the text that surrounds the hyperlink pointing to the page. Many researches (Chakrabarti et al., 1999;

Haveliwala et al., 2002; Varlamis et al., 2004; Gibson et al., 1998; Kumar et al., 1999) have done and solutions have suggested to the problem of searching, indexing or querying the Web, taking into account its structure as well as the meta-information included in the hyperlinks and the text surrounding them. There are a number of algorithms proposed based on the link analysis. Using citation analysis, co-citation algorithm (Dean and Henzinger, 1999) and extended co-citation algorithm (Hou and Zhang, 2003) are proposed. These algorithms are simple and deeper relationships among the pages cannot be discovered. Three important algorithms Page Rank (Brin and Page, 1998), Weighted Page Rank (WPR) (Xing and Ghorbani, 2004) and Hypertext Induced Topic Search (HITS) (Kleinberg, 1999a) are discussed below in detail and compared.

VI. APPLICATIONS

- a. Visually impaired can browse the net.
- b. To find the synonyms for technical words.
- c. Agriculture and Horticulture
- d. Court reporting
- e. Medicine
- f. Military
- g. Training Air-traffic Controllers
- h. Telephony and other domains
- i. Transcription

VII. CONCLUSION

Speech Recognition is widely used in industrial software market [20]. Researchers, working on the very promising and challenging field of Speech Recognition are bearing towards the ultimate goal i.e., Natural Conversation between Human beings and machines, are applying the knowledge from areas of Acoustic-Phonetics, Speech Perception, Artificial Intelligence etc., The challenges to the recognition performance of SR are being provided concrete solution so that the gap between recognition capability of machine and that a human being can be reduced to maximum extend [7]. Main goal of SR is to design a network which would be able to do continuous speech recognition on a larger vocabulary [20]. An attempt has been made through this paper give a comprehensive survey and uses of speech recognition. Speech Recognition and Speech Synthesis is mainly focused on Blind or visually impaired and Deaf or Hearing impaired. We focus on visually impaired persons to access the web without any problem. Visually Impaired persons can access the web through this Speech based information Retrieval through net. The required information can be obtained by the Audio query and this is converted into text and the relevant information can be retrieved and finally it has to be converted into Audio.

VIII. REFERENCES

- [1]. Aggarwal, R.K. and Dave, M., “Acoustic Modeling Problem for Automatic Speech Recognition System: Conventional Methods (Part I)”, International Journal of Speech Technology (2011) 14:297–308.
- [2]. Song Yang, Meng Joo Er, and Yang Gao. "A High Performance Neural-Networks-Based Speech Recognition System", IEEE trans.pp.1527, 2001.
- [3]. Kimberlee A. Kemble “An Introduction to Speech Recognition”,Voice Systems Middleware Education IBM Corporation.
- [4]. Digital Signal Processing Mini Project.” An Automatic Speaker Recognition System”, 14 June 2005.
- [5]. Samoulian, A., “Knowledge Based Approach to Speech Recognition”, 1994.
- [6]. “ECE3Speech Recognition.” A Simple Speech Recognition Algorithm.15 April 2003. 1 July 2005.
- [7]. Qiguo DUAN, Duoqian MIAO, ZHANG, Ji ZHENG, “Personalized web retrieval based on Rough-Fuzzy Method” Journals of Computational Information System(2007)203-208.
- [8]. P. Ravi Kumar and Ashutosh Kumar Singh, “Web Structure Mining: Exploring Hyperlinks and Algorithms for Information Retrieval”, American Journal of Applied Sciences 7 (6): 840-845, 2010. (Article in Journal).

- [9]. "Speech Recognition by Dynamic Time Warping", 20th April 1998. 06 July 2005.
- [10]. Corneliu Octavian DUMITRU, Inge GAVAT. "Vowel, Digit and Continuous Speech Recognition Based on Statistical, Neural and Hybrid Modeling by Using ASRS_RL ". EUROCON 2007, The International Conference on "Computer as Tool", pp.858-859. (Article in Conference Proceedings)
- [11]. Gavath, O.Dumgitru, C. Iancu, Gostache, "Learning strategies in speech Recognition", Proc. Elmar 2005, pp.237-240, June 2005, Zadar, Croatia.
- [12]. J. Bilmes: "A Gentle Tutorial on the EM Algorithm and its Application to Parameter Estimation for Gaussian Mixture and Hidden Markov Models", Technical Report, University of Berkeley, ICSI-TR-97-021. April 1998. 01 August 2005.
- [13]. Bahlmann. Haasdonk. Burkhardt. "Speech and audio recognition" IEEE trans. Vol. 11. May 2003.
- [14]. Edward Gatt, Joseph Micallef, Paul Micsllef, Edward Chilton. "Phoneme Classification in Hardware Implemented Neural Networks "IEEE trans, pp.481, 2001
- [15]. "Speech Recognition Based on Statistical, Neural and Hybrid Modeling by Using ASRS_RL ". EUROCON 2007, The International Conference on "Computer as Tool", pp.858-859.
- [16]. Leitch D., MacMillan T., Year III Final Research Report on the Liberated Learning Project "How Students with Disabilities Respond to Speech Recognition Technology in the University Classroom", July 2002.
- [17]. Language: Implications for Deaf Readers". Journal of Deaf Studies and Deaf Education 5(1). Winter 2000. 32-50.
- [18]. Dutch Television Programs: A Text Linguistic Approach". Journal of Deaf Studies and Deaf Education 10(4). Fall 2005. 402-416.
- [19]. Meysam Mohamad pour, Fardad Farokhi, "An Advanced method for Speech Recognition" , World Academy of science, Engineering and Technology 49 2009.

Short Bio Data for the Author

S. Ananthi received her B.E Degree in Computer Science from Dhanalakshmi Srinivasan Engineering College,

Perambalur and M.E Degree in Computer Science from Annamalai University, Chidambaram. She worked as an Assistant Professor in M.A.M College of Engineering, Siruganur, Trichy. Now currently she is pursuing her Doctorate in Computer Science and Engineering, Annamalai University, Chidambaram. She had presented a paper about Steganography in National Conference on Multimedia Signal Processing, NCMSP' 2011 in Annamalai University on 2011. Her Paper got Selected in the National Conference on Information Security (NCIS-2012) held at SASTRA University, Thanjavur, during June 14-15, 2012. She participated in the National Conference on Confluence of Multidisciplinary in Engineering Fields NCCME'12. She had published a paper in International Journal of Engineering and Innovative Technology (IJEIT). She had published a paper in International Organization of Scientific Research (IOSR) International Journal on the topic of Steganography. Her research interests include speech processing, cryptography and steganography.

Dr. P. Dhanalakshmi received her Bachelor's degree in Computer Science and Engineering from Government College of Technology, Coimbatore in the year 1993. She received her M.Tech degree in Computer Applications from the reputed Indian Institute of Technology, New Delhi under the Quality Improvement Programme in the year 2003. She completed her Ph. D in Computer Science and Engineering from Annamalai University in the year 2011. She joined the services of Annamalai University in the year 1998 as a faculty member and is presently serving as Associate Professor in the Department of Computer Science & Engg. She has published 11 papers in international conferences and journals. She is guiding several students who are pursuing doctoral research. Her research interests include speech processing, image and video processing, pattern classification and neural networks.