

論文 / 著書情報
Article / Book Information

Title	Recent Advances in Speaker Recognition (Invited paper)
Author	SADAOKI FURUI
Journal/Book name	Audio- and Video-based Biometric Person Authentication: First International Conference (AVBPA '97), Vol. , No. , pp. 237-252
発行日 / Issue date	1997, 4
DOI	http://dx.doi.org/10.1007/BFb0016001
権利情報 / Copyright	The original publication is available at www.springerlink.com .

Recent Advances in Speaker Recognition

Sadaoki Furui

NTT Human Interface Laboratories
3-9-11, Midori-cho, Musashino-shi, Tokyo, 180 Japan
Tokyo Institute of Technology
2-12-1, O-okayama, Meguro-ku, Tokyo, 152 Japan
furui@splab.hil.ntt.co.jp, furui@cs.titech.ac.jp

Abstract. This paper introduces recent advances in speaker recognition technology. The first part discusses general topics and issues. The second part is devoted to a discussion of more specific topics of recent interest that have led to interesting new approaches and techniques. They include VQ- and ergodic-HMM-based text-independent recognition methods, a text-prompted recognition method, parameter/distance normalization and model adaptation techniques, and methods of updating models and *a priori* thresholds in speaker verification. Although many recent advances and successes have been achieved in speaker recognition, there are still many problems for which good solutions remain to be found. The last part of this paper describes 16 open questions about speaker recognition. The paper concludes with a short discussion assessing the current status and future possibilities.

1. Principles of Speaker Recognition

1.1 What is Speaker Recognition?

Speaker recognition is the process of automatically recognizing who is speaking by using speaker-specific information included in speech waves [5][10][11][13][14][40][49]. This technique can be used to verify the identity claimed by people accessing systems; that is, it enables access control of various services by voice. Applicable services include voice dialing, banking over a telephone network, telephone shopping, database access services, information and reservation services, voice mail, security control for confidential information, and remote access of computers. Another important application of speaker recognition technology is its use for forensic purposes [23].

1.2 Classification of Speaker Recognition

Speaker recognition can be classified into speaker identification and speaker verification. Speaker identification is the process of determining from which of the registered speakers a given utterance comes. Speaker verification is the process of accepting or rejecting the identity claim of a speaker. Most of the applications in which voice is used to confirm the identity claim of a speaker are classified as speaker verification. The fundamental difference between identification and verification is the number of decision alternatives. In identification, the number of decision alternatives is equal to the size of the population, whereas in verification there are only two choices, accept or reject, regardless of the population size. Therefore, speaker identification performance decreases as the size of the population increases, whereas speaker verification performance approaches a constant, independent of the size of the population, unless the distribution of physical characteristics of speakers is extremely biased.

There is also the case called "open set" identification, in which a reference model for the unknown speaker may not exist. In this case, an additional decision alternative, "the unknown does not match any of the models", is required. In either verification or identification, an additional threshold test can be applied to determine whether the match is close enough to accept the decision or ask for a new trial.

Speaker recognition methods can also be divided into text-dependent and text-independent methods. The former require the speaker to provide utterances of key words or sentences that are the same text for both training and recognition, whereas the latter do

not rely on a specific text being spoken. The text-dependent methods are usually based on template-matching techniques in which the time axes of an input speech sample and each reference template or reference model of the registered speakers are aligned, and the similarity between them is accumulated from the beginning to the end of the utterance [9][37][48][60]. Since this method can directly exploit voice individuality associated with each phoneme or syllable, it generally achieves higher recognition performance than the text-independent method.

However, there are several applications, such as forensic and surveillance applications, in which predetermined key words cannot be used. Moreover, human beings can recognize speakers irrespective of the content of the utterance. Therefore, text-independent methods have recently attracted more attention. Another advantage of text-independent recognition is that it can be done sequentially, until a desired significance level is reached, without the annoyance of the speaker having to repeat the key words again and again.

Both text-dependent and independent methods have a serious weakness. That is, these systems can easily be defeated, because someone who plays back the recorded voice of a registered speaker uttering key words or sentences into the microphone can be accepted as the registered speaker. To cope with this problem, some methods use a small set of words, such as digits, as key words, and each user is prompted to utter a given sequence of key words that is randomly chosen every time the system is used [21][48]. Yet even this method is not reliable enough, since it can be defeated with advanced electronic recording equipment that can reproduce key words in a requested order. Therefore, a text-prompted speaker recognition method has recently been proposed. (See Section 2.3)

1.3 Basic Structures of Speaker Recognition Systems

In the speaker identification task, a speech utterance from an unknown speaker is analyzed and compared with speech models of known speakers. The unknown speaker is identified as the speaker whose model best matches the input utterance. In speaker verification, an identity claim is made by an unknown speaker, and an utterance of this unknown speaker is compared with the model for the speaker whose identity is claimed. If the match is good enough, that is, above a threshold, the identity claim is accepted. A high threshold makes it difficult for impostors to be accepted by the system, but at the risk of falsely rejecting valid users. Conversely, a low threshold enables valid users to be accepted consistently, but at the risk of accepting impostors. To set the threshold at the desired level of customer rejection and impostor acceptance, it is necessary to know the distribution of customer and impostor scores.

The effectiveness of speaker-verification systems can be evaluated by using the receiver operating characteristics (ROC) curve adopted from psychophysics. The ROC curve is obtained by assigning two probabilities, the probability of correct acceptance and the probability of incorrect acceptance, to the vertical and horizontal axes respectively, and varying the decision threshold [11]. The equal-error rate (EER) is a commonly accepted overall measure of system performance. It corresponds to the threshold at which the false acceptance rate is equal to the false rejection rate.

1.4 Feature Parameters

Speaker identity is correlated with the physiological and behavioral characteristics of the speech production system for each speaker. These characteristics exist both in the spectral envelope (vocal tract characteristics) and in the supra-segmental features (voice source characteristics) of speech. Although it is impossible to separate these kinds of characteristics, and many voice characteristics are difficult to measure explicitly, many characteristics are captured implicitly by various signal measurements. Signal measurements such as short-term and long-term spectra and overall energy are easy to obtain. These measurements provide the means for effectively discriminating among speakers. Fundamental frequency can also be used to recognize speakers if it can be extracted reliably [1][4][28].

Currently, the most commonly used short-term spectral measurements are LPC-derived cepstral coefficients and their regression coefficients [9]. A spectral envelope reconstructed from a truncated set of cepstral coefficients is much smoother than one reconstructed from LPC coefficients, and hence provides a stabler representation from one repetition to another of a particular speaker's utterances. As for the regression coefficients, typically, the first- and second-order coefficients, that is, derivatives of the time functions of cepstral coefficients, are extracted at every frame period to represent spectral dynamics. They are respectively called the delta-cepstral and delta-delta-cepstral coefficients.

1.5 Text-Dependent Speaker Recognition Methods

Text-dependent speaker recognition methods can be classified into DTW (dynamic time warping) or HMM (hidden Markov model) based methods.

1.5.1 DTW-based methods

A typical approach to text-dependent speaker recognition is the spectral template matching approach [9]. In this approach, each utterance is represented by a sequence of feature vectors, generally, short-term spectral feature vectors, and the trial-to-trial timing variation of utterances of the same text is normalized by aligning the analyzed feature vector sequence of a test utterance to the template feature vector sequence using a DTW algorithm. The overall distance between the test utterance and the template is used for recognition decision.

1.5.2 HMM-based methods

An HMM can efficiently model the statistical variation in spectral features. Therefore, HMM-based methods have achieved significantly better recognition accuracies than DTW-based methods [37][48][60].

A speaker verification system based on characterizing the utterances as sequences of subword units represented by HMMs has been introduced and tested [46]. Two types of subword units, phone-like units (PLUs) and acoustic segment units (ASUs), have been studied. PLUs are based on phonetic transcriptions of spoken utterances and ASUs are extracted directly from the acoustic signal without using any linguistic knowledge. The results of experiments using isolated digit utterances show only small differences in performance between PLU- and ASU-based representations.

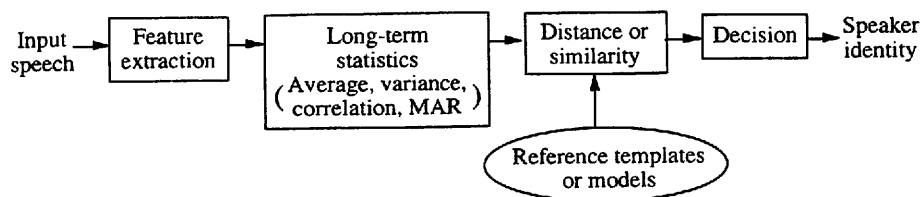
1.6 Text-Independent Speaker Recognition Methods

In text-independent speaker recognition, the words or sentences used in recognition trials generally cannot be predicted. Since it is impossible to model or match speech events at the word or sentence level, the following three kinds of methods shown in Fig. 1 have been actively investigated [10].

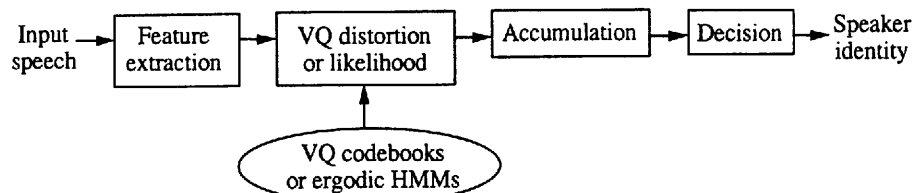
1.6.1 Long-term-statistics-based methods

As text-independent features, long-term sample statistics of various spectral features, such as the mean and variance of spectral features over a series of utterances, have been used [7][25][26] (Fig. 1(a)). However, long-term spectral averages are extreme condensations of the spectral characteristics of a speaker's utterances and, as such, lack the discriminating power included in the sequences of short-term spectral features used as models in text-dependent methods. In one of the trials using the long-term averaged spectrum [7], the effect of session-to-session variability was reduced by introducing a weighted cepstral distance measure.

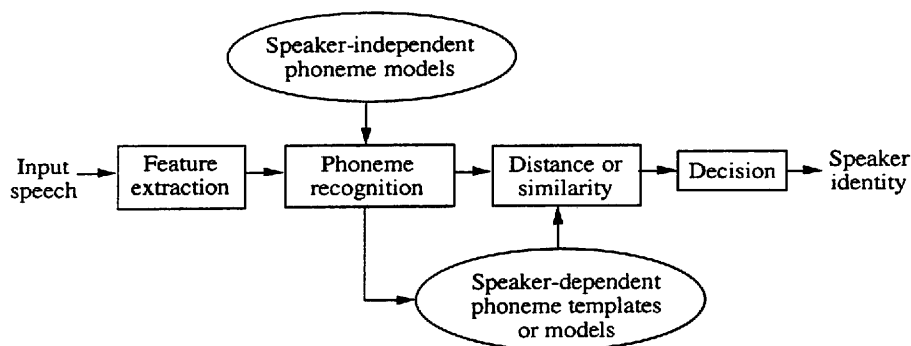
Studies on using statistical dynamic features have also been reported. Montacie et al. [36] applied a multivariate auto-regression (MAR) model to the time series of cepstral vectors to characterize speakers, and reported good speaker recognition results. Griffin et al. [19] studied distance measures for the MAR-based method, and reported that when ten sentences were used for training and one sentence was used for testing, identification and verification rates were almost the same as those obtained by an HMM-based method.



(a) Long-term-statistics-based method



(b) VQ / HMM-based method



(c) Speech-recognition-based method

Fig. 1. Basic structures of text-independent speaker recognition methods.

It was also reported that the optimum order of the MAR model was 2 or 3, and that distance normalization using *a posteriori* probability was essential to obtain good results in speaker verification.

1.6.2 VQ-based methods

A set of short-term training feature vectors of a speaker can be used directly to represent the essential characteristics of that speaker. However, such a direct representation is impractical when the number of training vectors is large, since the memory and amount of computation required become prohibitively large. Therefore, attempts have been made to find efficient ways of compressing the training data using vector quantization (VQ) techniques.

In this method, VQ codebooks, consisting of a small number of representative feature vectors, are used as an efficient means of characterizing speaker-specific features [24][28][29][45][53][55]. A speaker-specific codebook is generated by clustering the training feature vectors of each speaker. In the recognition stage, an input utterance is vector-quantized by using the codebook of each reference speaker; the VQ distortion accumulated over the entire input utterance is used for making the recognition determination (Fig. 1(b)).

1.6.3 Ergodic-HMM-based methods

The basic structure is the same as the VQ-based method (Fig. 1(b)), but in this method an ergodic HMM is used instead of a VQ codebook. Over a long timescale, the temporal variation in speech signal parameters is represented by stochastic Markovian transitions between states. Poritz [41] proposed using a five-state ergodic HMM (i.e., all possible transitions between states are allowed) to classify speech segments into one of the broad phonetic categories corresponding to the HMM states. A linear predictive HMM was adopted to characterize the output probability function. He characterized the automatically obtained categories as strong voicing, silence, nasal/liquid, stop burst/post silence, and frication.

Tishby [58] extended Poritz's work to the richer class of mixture autoregressive (AR) HMMs. In these models, the states are described as a linear combination (mixture) of AR sources. It can be shown that mixture models are equivalent to a larger HMM with simple states, together with additional constraints on the possible transitions between states.

1.6.4 Speech-recognition-based methods

The VQ- and HMM-based methods can be regarded as methods that use phoneme-class-dependent speaker characteristics in short-term spectral features through implicit phoneme-class recognition. In other words, phoneme-classes and speakers are simultaneously recognized in these methods. On the other hand, in the speech-recognition-based methods (Fig. 1(c)), phonemes or phoneme-classes are explicitly recognized, and then each phoneme (-class) segment in the input speech is compared with speaker models or templates corresponding to that phoneme (-class).

Savic et al. [51] used a five-state ergodic linear predictive HMM for broad phonetic categorization. In their method, after frames that belong to particular phonetic categories have been identified, feature selection is performed. In the training phase, reference templates are generated and verification thresholds are computed for each phonetic category. In the verification phase, after phonetic categorization, a comparison with the reference template for each particular category provides a verification score for that category. The final verification score is a weighted linear combination of the scores for each category. The weights are chosen to reflect the effectiveness of particular categories of phonemes in discriminating between speakers and are adjusted to maximize the verification performance. Experimental results showed that verification accuracy can be considerably improved by this category-dependent weighted linear combination method. Broad phonetic categorization can also be implemented by a speaker-specific hierarchical classifier instead of by an HMM, and the effectiveness of this approach has also been confirmed [6].

Recently, speaker verification through large vocabulary continuous speech recognition has been investigated. Details are given in Section 2.2.

2. Recent Advances

2.1 VQ- and HMM-Based Text-Independent Methods

Matsui et al. [28][29] tried a method using a VQ-codebook for long feature vectors consisting of instantaneous and transitional features calculated for both cepstral coefficients and fundamental frequency. Since the fundamental frequency cannot be

extracted from unvoiced speech, there are two separate codebooks for voiced and unvoiced speech for each speaker. A new distance measure was introduced to take into account the intra- and inter-speaker variability and to deal with the outlier problem in the distribution of feature vectors. The outlier vectors correspond to intersession spectral variation and to the difference between phonetic content of the training texts and the test utterances. Experimental results confirmed high recognition accuracies even when the codebooks for each speaker were made using training utterances recorded in a single session and the time difference between training and testing was more than three months. It was also confirmed that, although the fundamental frequency achieved only a low recognition rate by itself, the recognition accuracy was greatly improved by combining the fundamental frequency with spectral envelope features.

In contrast with the memoryless VQ-based method, non-memoryless source coding algorithms have also been studied using a segment (matrix) quantization technique [22][57]. The advantage of a segment quantization codebook over a VQ codebook representation is its characterization of the sequential nature of speech events. Higgins and Wohlford [20] proposed a segment modeling procedure for constructing a set of representative time normalized segments, which they called "filler templates". The procedure, a combination of K-means clustering and dynamic programming time alignment, provided a means for handling temporal variation.

Matsui et al. [30] compared the VQ-based method with the discrete/continuous ergodic HMM-based method, particularly from the viewpoint of robustness against utterance variations. They found that the continuous ergodic HMM method is far superior to the discrete ergodic HMM method and that the continuous ergodic HMM method is as robust as the VQ-based method when enough training data is available. However, when little data is available, the VQ-based method is more robust than the continuous HMM method. They investigated speaker identification rates using the continuous HMM as a function of the number of states and mixtures. It was shown that the speaker recognition rates are strongly correlated with the total number of mixtures, irrespective of the number of states. This means that the information on transitions between different states is ineffective for text-independent speaker recognition.

The case of a single-state continuous ergodic HMM corresponds to the technique based on the maximum likelihood estimation of a Gaussian-mixture model representation investigated by Rose et al. [43]. Furthermore, the VQ-based method can be regarded as a special (degenerate) case of a single-state HMM with a distortion measure being used as the observation probability. Gaussian mixtures are noted for their robustness as a parametric model and their ability to form smooth estimates of rather arbitrary underlying densities.

The ASU-based speaker verification method described in Section 1.5.2 has also been tested in the text-independent mode [47]. It has been shown that this approach can be extended to large vocabularies and continuous speech.

2.2 Speech-Recognition-Based Speaker Recognition

Gauvain et al. [16] investigated a statistical modeling approach, where each speaker was viewed as a source of phonemes, modeled by a fully connected Markov chain. Maximum *a posteriori* (MAP) estimation was used to generate speaker-specific models from a set of speaker-independent seed models. The lexical and syntactic structures of the language were approximated by local phonotactic constraints. The unknown speech is recognized by all of the speakers' models in parallel, and the hypothesized identity is that associated with the model set having the highest likelihood.

Since phonemes and speakers are simultaneously recognized by using speaker-specific Markov chains, this method can be considered as an extension of the ergodic-HMM-based method. The experimental results using the BREF corpus showed that this method clearly out-performed a simpler Gaussian mixture (single-state HMM) model. It was also found that text-independent and text-dependent verification EERs were about the same.

Rosenberg et al. have been testing a speaker verification system using 4-digit phrases

under field conditions of a banking application [48][52]. In this system, input speech is segmented into individual digits using a speaker-independent HMM. The frames within the word boundaries for a digit are compared with the corresponding speaker-specific HMM digit model and the Viterbi likelihood score is computed. This is done for each of the digits making up the input utterance. The verification score is defined to be the average normalized log-likelihood score over all the digits in the utterance.

Newman et al. [39] used a large vocabulary speech recognition system for speaker verification. A set of speaker-independent phoneme models were adapted to each speaker. The speaker verification consisted of two stages. First, speaker-independent speech recognition was run on each of the test utterances to obtain phoneme segmentation. In the second stage, the segments were scored against the adapted models for a particular target speaker. The scores were normalized by those with speaker-independent models. The system was evaluated using the 1995 NIST-administered speaker verification database, which consists of data taken from the Switchboard corpus. The results showed that this method could not out-perform Gaussian mixture models.

2.3 Text-Prompted Speaker Recognition

How can we prevent speaker verification systems from being defeated by a recorded voice? Another problem is that people often do not like text-dependent systems because they do not like to utter their identification number, such as their social security number, within the hearing of other people. To cope with these problems, a text-prompted speaker recognition method has been proposed.

In this method, key sentences are completely changed every time [31][33]. The system accepts the input utterance only when it determines that the registered speaker uttered the prompted sentence. Because the vocabulary is unlimited, prospective impostors cannot know in advance the sentence they will be prompted to say. This method can not only accurately recognize speakers, but can also reject an utterance whose text differs from the prompted text, even if it is uttered by a registered speaker. Thus, a recorded and played back voice can be correctly rejected.

This method uses speaker-specific phoneme models as basic acoustic units. One of the major issues in this method is how to properly create these speaker-specific phoneme models when using training utterances of a limited size. The phoneme models are represented by Gaussian-mixture continuous HMMs or tied-mixture HMMs, and they are made by adapting speaker-independent phoneme models to each speaker's voice. Since the text of training utterances is known, these utterances can be modeled as the concatenation of phoneme models, and these models can be automatically adapted by an iterative algorithm.

In the recognition stage, the system concatenates the phoneme models of each registered speaker to create a sentence HMM, according to the prompted text. Then the likelihood of input speech against the sentence model is calculated and used for the speaker recognition determination. If the likelihood of both speaker and text is high enough, the speaker is accepted as the claimed speaker. Experimental results gave a speaker and text verification rate of 99.4% when the adaptation method for tied-mixture-based phoneme models and the likelihood normalization method described in the Section 2.4.2 were used.

2.4 Normalization and Adaptation Techniques

How can we normalize the intra-speaker variation of likelihood (similarity) values in speaker verification? The most significant factor affecting automatic speaker recognition performance is variation in signal characteristics from trial to trial (intersession variability or variability over time). Variations arise from the speaker him/herself, from differences in recording and transmission conditions, and from noise. Speakers cannot repeat an utterance precisely the same way from trial to trial. It is well known that samples of the same utterance recorded in one session are much more highly correlated than tokens recorded in separate sessions. There are also long term trends in voices [7][8].

It is important for speaker recognition systems to accommodate these variations.

Adaptation of the reference model as well as the verification threshold for each speaker is indispensable to maintain a high recognition accuracy for a long period. In order to compensate for the variations, two types of normalization techniques have been tried — one in the parameter domain, and the other in the distance/similarity domain. The latter technique uses the likelihood ratio or a *posteriori* probability. To adapt HMMs for noisy conditions, the HMM composition (PMC: parallel model combination) method has been successfully tried.

2.4.1 Parameter-domain normalization

As one typical normalization technique in the parameter domain, spectral equalization, the so-called “blind equalization” method, has been confirmed to be effective in reducing linear channel effects and long-term spectral variation [2][9]. This method is especially effective for text-dependent speaker recognition applications using sufficiently long utterances. In this method, cepstral coefficients are averaged over the duration of an entire utterance, and the averaged values are subtracted from the cepstral coefficients of each frame. This method can compensate fairly well for additive variation in the log spectral domain. However, it unavoidably removes some text-dependent and speaker-specific features, so it is inappropriate for short utterances in speaker recognition applications.

It was shown that time derivatives of cepstral coefficients (delta-cepstral coefficients) are resistant to linear channel mismatch between training and testing [9][56].

2.4.2 Likelihood normalization

Higgins et al. [21] proposed a normalization method for distance (similarity or likelihood) values that uses a likelihood ratio:

$$\log L(X) = \log p(X | S = S_c) - \log p(X | S \neq S_c) \quad (1)$$

The likelihood ratio is the ratio of the conditional probability of the observed measurements of the utterance given the claimed identity is correct to the conditional probability of the observed measurements given the speaker is an impostor. Generally, a positive value of $\log L$ indicates a valid claim, whereas a negative value indicates an impostor. We call the second term on the right hand side of Eq. (1) the normalization term.

The density at point X for all speakers other than true speaker S can be dominated by the density for the nearest reference speaker, if we assume that the set of reference speakers is representative of all speakers. This means that likelihood ratio normalization approximates optimal scoring in Bayes' sense. This normalization method is, however, unrealistic because, even if only the nearest reference speaker is used, conditional probabilities must be calculated for all the reference speakers, which costs a lot. Therefore, a set of speakers, “cohort speakers”, has been chosen for calculating the normalization term of Eq. (1). Higgins et al. proposed using speakers that are representative of the population near the claimed speaker.

An experiment in which the size of the cohort speaker set was varied from 1 to 5 showed that speaker verification performance increases as a function of the cohort size, and that the use of normalization significantly compensates for the degradation obtained by comparing verification utterances recorded using an electret microphone with models constructed from training utterances recorded with a carbon button microphone [50].

Matsui et al. [31][32] proposed a normalization method based on a *posteriori* probability:

$$\log L(X) = \log p(X | S = S_c) - \log \sum_{S \in \text{Ref}} p(X | S) \quad (2)$$

The difference between the normalization method based on the likelihood ratio and that based on a *posteriori* probability is whether or not the claimed speaker is included in the impostor speaker set for normalization; the cohort speaker set in the likelihood-ratio-

based method does not include the claimed speaker, whereas the normalization term for the *a posteriori*-probability-based method is calculated by using a set of speakers including the claimed speaker. Experimental results indicate that both normalization methods almost equally improve speaker separability and reduce the need for speaker-dependent or text-dependent thresholding, compared with scoring using only the model of the claimed speaker [32][50].

The normalization method using the cohort speakers that are representative of the population near the claimed speaker is expected to increase the selectivity of the algorithm against voices similar to the claimed speaker. However, this method has a serious problem that it is vulnerable to attack by impostors of the opposite gender. Since the cohorts generally model only same-gender speakers, the probability of opposite-gender impostor speech is not well modeled and the likelihood ratio is based on the tails of distributions, which gives rise to unreliable values. Another way of choosing the cohort speaker set is to use speakers who are typical of the general population. Reynolds [42] reported that a randomly selected, gender-balanced background speaker population outperformed a population near the claimed speaker.

Carey et al. [3] proposed a method in which the normalization term is approximated by the likelihood for a world model representing the population in general. This method has an advantage that the computational cost for calculating the normalization term is much smaller than the original method since it does not need to sum the likelihood values for cohort speakers. Matsui et al. [32] recently proposed a new method based on tied-mixture HMMs in which the world model is made as a pooled mixture model representing the parameter distribution for all the registered speakers. This model is created by averaging together the mixture-weighting factors of each reference speaker calculated using speaker-independent mixture distributions. Therefore the pooled model can be easily updated when a new speaker is added as a reference speaker. In addition, this method has been confirmed to give much better results than either of the original normalization methods.

Since these normalization methods neglect the absolute deviation between the claimed speaker's model and the input speech, they cannot differentiate highly dissimilar speakers. Higgins et al. [21] reported that a multilayer network decision algorithm makes effective use of the relative and absolute scores obtained from the matching algorithm.

2.4.3 HMM adaptation for noisy conditions

How can we improve the robustness of speaker recognition techniques against noisy speech or speech distorted by a telephone? The robustness issue is crucial in real-world applications. Will speaker recognition technology ever be capable of performing with high reliability under adverse real-world conditions (noisy telephone lines, limited training data, etc.)? Only a few papers have addressed this problem in speaker recognition.

Rose et al. [44] applied the HMM composition (PMC: parallel model combination) method [15][27] to speaker identification under noisy conditions. The HMM composition is a technique to combine a clean speech HMM and a background noise HMM to create a noise-added speech HMM. They can be optimally combined by using the expectation-maximization (EM) algorithm. Experimental results show that this method is highly effective at recognizing speech with additive noise. In this method, the signal-to-noise ratio (SNR) of input speech is assumed to be similar to that of training speech or to be given. However, the SNR is variable in real situations and it is also difficult to measure especially when the noise is non-stationary. Matsui et al. [35] proposed a method in which several noise-added HMMs with various SNRs were created and the HMM that had the highest likelihood value for the input speech was selected. A speaker decision was made using the likelihood value corresponding to the selected model. Experimental application of this method to text-independent speaker identification and verification in various kinds of noisy environments demonstrated considerable improvement in speaker recognition.

2.5 Updating Models and A Priori Threshold for Speaker Verification

How do we deal with long-term variability in people's voices? How should we update the speaker models to cope with the gradual changes in people's voices? Since we cannot ask every user to utter many utterances at many different sessions in real situations, it is necessary to build each speaker model based on a small amount of data collected at a few sessions, and then the model must be updated using speech data collected when the system is used. How do we adequately retrain the models?

How should we set the a priori decision threshold for speaker verification? In most laboratory speaker recognition experiments, the threshold is set *a posteriori* so that the equal error rate (EER) is achieved. Since the threshold cannot be set *a posteriori* in real situations, we have to have reasonable ways to set the threshold before verification. It must be set according to the importance of the two errors, which depends on the application.

These two problems are highly related each other. Furui [9] proposed methods for updating reference templates and the threshold in DTW-based speaker verification. An optimum threshold was estimated based on the distribution of overall distances between each speaker's reference template and a set of utterances of other speakers (interspeaker distances). The interspeaker distance distribution was approximated by a normal distribution, and the threshold was calculated by the linear combination of its mean value and standard deviation. The intraspeaker distance distribution was not taken into account in the calculation, mainly because it is difficult to obtain stable estimates of the intraspeaker distance distribution for small numbers of training utterances. It is easy to estimate interspeaker distance distributions by cross comparison of the training utterances between different speakers. The threshold was updated at the same time as the reference templates. The reference template for each speaker was updated by averaging new utterances and the present template after time registration.

Matsui et al. [34] extended these methods and applied them to text-independent and text-prompted speaker verification using HMMs. HMM parameters for each speaker were updated by using present parameters and a limited amount of data for adaptation so that they were close to the parameter values given by maximum likelihood estimation using all the past and present data. The thresholds for speaker verification were updated according to Eq. (3). The threshold converges from a threshold value that has a higher false acceptance (FA) rate to the EER threshold value for the samples for updating the model as the updating of the model proceeds.

$$\bar{\phi} = \omega\phi_1 + (1 - \omega)\phi_0 \quad (3)$$

$$\omega = \frac{2}{1 + \exp(a \cdot k)} \quad (4)$$

Here, ϕ_0 is the e.e.t. for the samples for updating the model; ϕ_1 is the threshold which determines the upper bound of the FA rate (Fig. 2). In the experiment, ϕ_1 was set to where the FA rate was slightly higher than the EER. The w is a parameter for controlling the convergence of the threshold, and for instance, defined as (4) with k being the number of sessions for updating the model and a being an experimental parameter. Evaluation of the performance of the two methods using 20 male speakers showed that the verification error rate was about 40% of that without updating in text-independent experiments, and about 80% in text-prompted ones.

Setlur et al. [52] showed that speaker verification performance was improved by constraining the individual feature variances for all models to a fixed set of values when the size of training data was small. This is because the model variance estimates derived from the K-means clustering approach using small data are poor and result in erratic performance. The fixed set of variances was derived by averaging the model variances of the speaker-independent model across all states and mixtures for each feature. The decision threshold was determined *a priori* by using an independent impostor set and a target false acceptance rate.

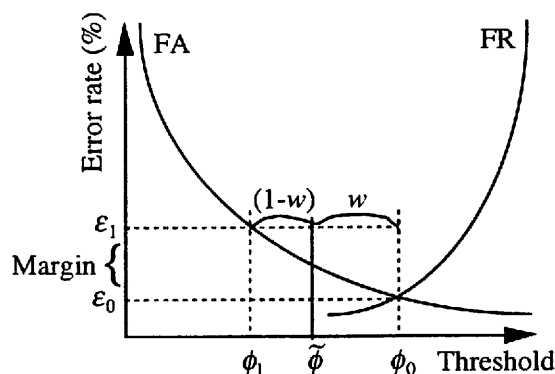


Fig. 2. Method of updating the speaker verification threshold.

3. Open Questions about Speaker Recognition

Although many recent advances and successes have been achieved as described in the previous sections, there are still many problems for which good solutions remain to be found. Sixteen major problems are discussed below.

(1) How can human beings correctly recognize speakers?

Research on perceptual voice individuality has been conducted for many years, but what we have learned about the hearing capability of human beings is very limited. Which acoustic parameters do listeners utilize in making their judgments? One of the results found by experiment is that segmental (spectral envelope) information plays a more important role than supra-segmental (pitch and energy) information. It is very difficult to give an answer to a question like "which performs better, a human being or a computer". The answer very much depends on the conditions, such as the number of speakers, their familiarity to the listeners, noise and distortion, and the time difference between training and testing.

(2) Is it useful to study the mechanism of speaker recognition by human beings?

Do auditory models help improve speaker recognition performance? An example of how knowledge of human hearing characteristics led to a new idea for automatic speaker recognition was the creation of delta-cepstral parameters (polynomial expansion coefficients of cepstral parameters). These parameters were proposed based on the experimental results that our hearing systems are very sensitive to spectral transition. Although the number of such examples is very small, we expect advancements in our understanding of the mechanism of speech perception to help create new engineering ideas in the future.

(3) Is it useful to study the physiological mechanism of speech production to get new ideas for speaker recognition?

How is identity encoded and is this worth pursuing for automatic speaker recognition? How are perceptually identical voices produced by our vocal systems?

(4) What feature parameters are appropriate for speaker recognition?

The most common parameters we use are cepstral and delta-cepstral parameters. These are exactly the same as those used in speech recognition. Is it worth searching for a new set of feature parameters that is more appropriate to speaker recognition than the present set?

(5) How can we model the relatively macro-transitional features covering the interval between 200 to 300 ms?

The delta-cepstral parameters usually represent the transitional features covering the interval between 50 to 100 ms. The macro-transitional features covering the interval between 200 to 300 ms correspond to phoneme-to-phoneme and syllable-to-syllable transitions. We do not know how to represent these transitional features for speaker recognition. As described in Section 2.1, it was found that transitional information between different states in HMM is ineffective for text-independent speaker recognition.

(6) How can we fully exploit the clearly evident encoding of identity in prosody and other suprasegmental features of speech?

Although it is quite obvious that suprasegmental features of speech play important roles in speech perception, it is still difficult to use these features in both speaker recognition and speech recognition. This is because (a) it is difficult to automatically extract these features correctly, (b) we do not know how to model these dynamic features, (c) they are highly variable, and (d) speakers can exert more conscious control over them than over segmental features.

(7) Is the "sheep and goats" problem (a small percentage of speakers account for the majority of errors) universal?

Is there any universal set of parameters, algorithms, etc. that is good for all/any set of speakers? Or is the encoding of identity so speaker-dependent that the optimal configuration parameters (algorithm, acoustic features, etc.) depend on the speaker set to be identified? Should we combine separate features to identify a wide range of speakers? Will speaker recognition ever be practical for 100% of the speaking population (or will some speakers need to be excluded because their "identity encoding scheme" cannot be exploited by the system configuration)? Is there any set of parameters that is good for separating speakers whose voices sound identical, such as twins? Can we separate these identically sounding voices? Is it worth trying?

(8) Can we ever reliably cluster speakers on the basis of similarity/dissimilarity?

Is there a "universal" (text independent, feature set independent, etc.) relationship between pairs of speakers? Can we make absolute statements about the relationships of speakers (e.g., these two speakers will never be confused)? We can find some systematic speaker variation across phonemes, but we do not know how systematic and how random it is. This question is very closely related to question (7).

(9) How do we acquire realistically sized (viable) databases that still adequately model inter- and intra-speaker variability?

How can we model or sample the universal distribution of voices (inter-speaker variability)? How should we choose the speaker set? How should we collect speech data for each speaker in order to capture the intra-speaker variability? How should we select the texts (sentences, words, syllables, etc.)? How should we make common speech databases for evaluating speaker recognition techniques? It is crucial to have common speech databases for comparing the effectiveness of different techniques [18][38]. Major speech databases designed for speaker recognition and related areas include the KING corpus and the SWITCHBOARD corpus [18]. It is also important to consider the recording conditions, such as background noise, the type of microphone/telephone, and channel characteristics, in order to collect a database appropriate for evaluating the robustness of the techniques.

(10) How do we deal with long-term variability in people's voices?

How should we update the speaker models to cope with the gradual changes in people's voices? Is there any systematic long-term variation? Can we estimate the session-to-session intra-speaker variability just by using speech samples collected at one session? How do we adequately retrain the models? How can we prevent the models from being

updated by the wrong persons' voices?

(11) Can we model or develop strategies for dealing with factors that significantly alter a person's voice (short-term transitory)?

Sources of these alterations include short-term illness (flu, etc.), emotion, fatigue, and the words spoken. So far, no one has succeeded in modeling these voice alterations.

(12) How can we extract text-independent speaker-specific features?

Various methods have been investigated in text-independent speaker recognition. However, there is still no good method for modeling transitional speaker characteristics. Another important issue is how to extract phoneme-dependent speaker characteristics without knowing the phonemes.

(13) How can we deal with deliberately disguised voices?

Are there acoustic features that are invariant regardless of the disguise method? This question is of particular importance in forensic speaker recognition. Criminals may disguise their voices or mimic another person's voice. Does this pose a real threat to speaker recognition systems? How can we cope with this problem?

(14) How does speaker recognition compare to other means of personal identification, both now and in the future?

Fingerprints, signatures, and various other means can also be used for identification. The cost/performance ratio of the speaker recognition system must become lower than that of the other means before it can be widely used in the real world. In contrast with finger-printing, we probably have to assume that there are pairs of speakers whose voices cannot be separated when the population is large.

(15) What are the conditions that speaker recognition systems must satisfy in order for them to be utilized in the field?

This question is a generalization of question (14). What is speech good for? Cost/performance is not necessarily the most important issue in actual use. Even if the performance is not very high, the system can be used in combination with other means. Although they are difficult to fully understand, the critical conditions that must be met by commercial systems are crucial.

(16) What about combining speech and speaker recognition?

The text-prompted method is an example of combining speech and speaker recognition. Since phonetic information and speaker information are correlated, combining the approaches and ideas of speaker and speech recognition is expected to create new ideas and improve recognition performance. But how can we do it? An interesting research topic is the automatic adjustment of speaker-independent phoneme models to each new speaker so that the performance of both speech and speaker recognition are simultaneously improved.

4. Concluding Remarks

There have been many recent advances and successes in speaker recognition technology. AT&T and TI (with Sprint) have started field tests and actual application of speaker recognition technology. However, there are still many problems for which good solutions remain to be found. Most of these problems arise from variability, including speaker-generated variability and variability in channel and recording conditions. It is very important to investigate feature parameters that are stable over a long period, insensitive to variations in speaking manner, including speaking rate and level, and robust against variations in voice quality such as those due to voice disguise or colds. It is also important to develop a method to cope with the problems of distortion due to telephone sets and channels and background and channel noises.

Speaker characterization techniques are related to research on improving speech recognition accuracy by speaker adaptation [12], improving synthesized speech quality by adding the natural characteristics of voice individuality, and converting synthesized voice individuality from one speaker to another. Studies on automatically extracting the speech periods of each person separately from a dialogue involving more than two people have recently appeared as an extension of speaker recognition technology [17][54][59]. Diversified research related to speaker-specific information in speech waves is expected to increase in the near future.

References

- [1] B. S. Atal, "Automatic Speaker Recognition Based on Pitch Contours," *J. Acoust. Soc. Am.*, Vol. 52, No. 6, pp. 1687-1697 (1972).
- [2] B. S. Atal, "Effectiveness of Linear Prediction Characteristics of the Speech Wave for Automatic Speaker Identification and Verification," *J. Acoust. Soc. Am.*, Vol. 55, No. 6, pp. 1304-1312 (1974).
- [3] M. J. Carey and E. S. Parris, "Speaker Verification Using Connected Words," *Proc. Institute of Acoustics*, Vol. 14, Part 6, pp. 95-100 (1992).
- [4] M. J. Carey, E. S. Parris, H. Lloyd-Thomas and S. Bennet, "Robust Prosodic Features for Speaker Identification," *Proc. Int. Conf. Spoken Language Processing*, Philadelphia, pp. 1800-1803 (1996).
- [5] G. R. Doddington, "Speaker Recognition—Identifying People by their Voices," *Proc. IEEE*, Vol. 73, No. 11, pp. 1651-1664 (1985).
- [6] J. Eatock and J. S. Mason, "Automatically Focusing on Good Discriminating Speech Segments in Speaker Recognition," *Proc. Int. Conf. Spoken Language Processing*, 5.2, pp. 133-136 (1990).
- [7] S. Furui, F. Itakura and S. Saito, "Talker Recognition by Longtime Averaged Speech Spectrum," *Trans. IECE*, 55-A, Vol. 1, No. 10, pp. 549-556 (1972).
- [8] S. Furui, "An Analysis of Long-Term Variation of Feature Parameters of Speech and its Application to Talker Recognition," *Trans. IECE*, 57-A, Vol. 12, pp. 880-887 (1974).
- [9] S. Furui, "Cepstral Analysis Technique for Automatic Speaker Verification," *IEEE Trans. Acoust. Speech, Signal Processing*, Vol. 29, No. 2, pp. 254-272 (1981).
- [10] S. Furui, "Research on Individuality Features in Speech Waves and Automatic Speaker Recognition Techniques," *Speech Communication*, Vol. 5, No. 2, pp. 183-197 (1986).
- [11] S. Furui, "Digital Speech Processing, Synthesis, and Recognition," Marcel Dekker, New York (1989).
- [12] S. Furui, "Speaker-Independent and Speaker-Adaptive Recognition Techniques," in *Advances in Speech Signal Processing* (eds. S. Furui and M. M. Sondhi), Marcel Dekker, New York, pp. 597-622 (1991).
- [13] S. Furui, "Speaker-Dependent-Feature Extraction, Recognition and Processing Techniques," *Speech Communication*, Vol. 10, No. 5-6, pp. 505-520 (1991).
- [14] S. Furui, "An Overview of Speaker Recognition Technology," *ESCA Workshop on Automatic Speaker Recognition, Identification and Verification*, pp. 1-9 (1994).
- [15] M. J. F. Gales and S. J. Young, "HMM Recognition in Noise Using Parallel Model Combination," *Proc. Eurospeech*, Berlin, pp. II-837-840 (1993).
- [16] J. L. Gauvain, L. F. Lamel and B. Proust, "Experiments with Speaker Verification over the Telephone," *Proc. Eurospeech*, Madrid, pp. 651-654 (1995).
- [17] H. Gish, M. -H. Siu and R. Rohlicek, "Segregation of Speakers for Speech Recognition and Speaker Identification," *Proc. IEEE Int. Conf. Acoust. Speech, Signal Processing*, Toronto, S13.11, pp. 873-876 (1991).
- [18] J. Godfrey, D. Graff and A. Martin, "Public Databases for Speaker Recognition and Verification," *ESCA Workshop on Automatic Speaker Recognition, Identification and Verification*, pp. 39-42 (1994).
- [19] C. Griffin, T. Matsui and S. Furui, "Distance Measures for Text-Independent Speaker Recognition Based on MAR Model," *Proc. IEEE Int. Conf. Acoust. Speech, Signal Processing*, Adelaide, 23.6, pp. I-309-312 (1994).

- [20] A. L. Higgins and R. E. Wohlford, "A New Method of Text-Independent Speaker Recognition," *Proc. IEEE Int. Conf. Acoust., Speech, Signal Processing*, 17.3, pp. 869-872 (1986).
- [21] A. Higgins, L. Bahler and J. Porter, "Speaker Verification Using Randomized Phrase Prompting," *Digital Signal Processing*, Vol. 1, pp. 89-106 (1991).
- [22] B. -H. Juang and F. K. Soong, "Speaker Recognition Based on Source Coding Approaches," *Proc. IEEE Int. Conf. Acoust., Speech, Signal Processing*, S5.4, pp. 613-616 (1990).
- [23] H. J. Kunzel, "Current Approaches to Forensic Speaker Recognition," *ESCA Workshop on Automatic Speaker Recognition, Identification and Verification*, pp.135-141 (1994).
- [24] K. -P. Li and E. H. Wrench Jr., "An Approach to Text-Independent Speaker Recognition with Short Utterances," *Proc. IEEE Int. Conf. Acoust., Speech, Signal Processing*, 12.9, pp. 555-558 (1983).
- [25] J. D. Markel, B. T. Oshika and A. H. Gray, "Long-Term Feature Averaging for Speaker Recognition," *IEEE Trans. Acoust. Speech Signal Processing*, Vol. ASSP-25, No. 4, pp. 330-337 (1977).
- [26] J. D. Markel and S. B. Davi, "Text-Independent Speaker Recognition from a Large Linguistically Unconstrained Time-Spaced Data Base," *IEEE Trans. Acoust. Speech Signal Processing*, Vol. ASSP-27, No. 1, pp. 74-82 (1979).
- [27] F. Martin, K. Shikano and Y. Minami, "Recognition of Noisy Speech by Composition of Hidden Markov Models," *Proc. Eurospeech*, Berlin, pp. II-1031-1034 (1993).
- [28] T. Matsui and S. Furui, "Text-Independent Speaker Recognition Using Vocal Tract and Pitch Information," *Proc. Int. Conf. Spoken Language Processing*, Kobe, 5.3, pp. 137-140 (1990).
- [29] T. Matsui and S. Furui, "A Text-Independent Speaker Recognition Method Robust Against Utterance Variations," *Proc. IEEE Int. Conf. Acoust. Speech Signal Processing*, S6.3, pp. 377-380 (1991).
- [30] T. Matsui and S. Furui, "Comparison of Text-Independent Speaker Recognition Methods Using VQ-Distortion and Discrete/Continuous HMMs," *Proc. IEEE Int. Conf. Acoust. Speech, Signal Processing*, San Francisco, pp. II-157-160 (1992).
- [31] T. Matsui and S. Furui, "Concatenated Phoneme Models for Text-Variable Speaker Recognition," *Proc. IEEE Int. Conf. Acoust. Speech, Signal Processing*, Minneapolis, pp. II-391-394 (1993).
- [32] T. Matsui and S. Furui, "Similarity Normalization Method for Speaker Verification Based on a Posteriori Probability," *ESCA Workshop on Automatic Speaker Recognition, Identification and Verification*, pp. 59-62 (1994).
- [33] T. Matsui and S. Furui, "Speaker Adaptation of Tied-Mixture-Based Phoneme Models for Text-Prompted Speaker Recognition," *Proc. IEEE Int. Conf. Acoust. Speech, Signal Processing*, Adelaide, 13.1 (1994).
- [34] T. Matsui and S. Furui, "Robust Methods of Updating Model and A Priori Threshold in Speaker Verification," *Proc. IEEE Int. Conf. Acoust. Speech, Signal Processing*, Atlanta, pp. I-97-100 (1996).
- [35] T. Matsui and S. Furui, "Speaker Recognition Using HMM Composition in Noisy Environments," *Computer Speech and Language*, Vol. 10, pp. 107-116 (1996).
- [36] C. Montacie et al., "Cinematic Techniques for Speech Processing: Temporal Decomposition and Multivariate Linear Prediction," *Proc. IEEE Int. Conf. Acoust. Speech, Signal Processing*, San Francisco, pp. I-153-156 (1992).
- [37] J. M. Naik, L. P. Netsch, and G. R. Doddington, "Speaker Verification over Long Distance Telephone Lines," *Proc. IEEE Int. Conf. Acoust., Speech, Signal Processing*, S10b.3, pp. 524-527 (1989).
- [38] J. Naik, "Speaker Verification over the Telephone Network: Databases, Algorithms and Performance Assessment," *ESCA Workshop on Automatic Speaker Recognition, Identification and Verification*, pp. 31-38 (1994).
- [39] M. Newman, L. Gillick, Y. Ito, D. McAllaster and B. Peskin, "Speaker Verification through Large Vocabulary Continuous Speech Recognition," *Proc. Int. Conf. Spoken Language Processing*, Philadelphia, pp. 2419-2422 (1996).
- [40] D. O' Shaughnessy, "Speaker Recognition," *IEEE ASSP Magazine*, 3, No. 4, pp. 4-17 (1986).
- [41] A. B. Poritz, "Linear Predictive Hidden Markov Models and the Speech Signal," *Proc. IEEE*

- Int. Conf. Acoust., Speech, Signal Processing, S11.5, pp. 1291-1294 (1982).
- [42] D. Reynolds, "Speaker Identification and Verification Using Gaussian Mixture Speaker Models," ESCA Workshop on Automatic Speaker Recognition, Identification and Verification, pp.27-30 (1994).
 - [43] R. C. Rose and R. A. Reynolds, "Text Independent Speaker Identification Using Automatic Acoustic Segmentation," Proc. IEEE Int. Conf. Acoust., Speech, Signal Processing, S51.10, pp. 293-296 (1990).
 - [44] R. C. Rose, E. M. Hofstetter, and D. A. Reynolds, "Integrated Models of Signal and Background with Application to Speaker Identification in Noise," IEEE Trans. Speech and Audio Processing, Vol. 2, No. 2, pp. 245-257 (1994).
 - [45] A. E. Rosenberg and F. K. Soong, "Evaluation of a Vector Quantization Talker Recognition System in Text Independent and Text Dependent Modes," Computer Speech and Language, 22, pp. 143-157 (1987).
 - [46] A. E. Rosenberg, C. -H. Lee and F. K. Soong, "Sub-Word Unit Talker Verification Using Hidden Markov Models," Proc. IEEE Int. Conf. Acoust., Speech, Signal Processing, S5.3, pp. 269-272 (1990).
 - [47] A. E. Rosenberg, C. -H. Lee, F. K. Soong and M. A. McGee, "Experiments in Automatic Talker Verification Using Sub-Word Unit Hidden Markov Models," Proc. Int. Conf. Spoken Language Processing, 5.4, pp. 141-144 (1990).
 - [48] A. E. Rosenberg, C. -H. Lee, and S. Gokcen, "Connected Word Talker Verification Using Whole Word Hidden Markov Models," Proc. IEEE Int. Conf. Acoust. Speech, Signal Processing, Toronto, S6.4, pp. 381-384 (1991).
 - [49] A. E. Rosenberg and F. K. Soong, "Recent Research in Automatic Speaker Recognition," in *Advances in Speech Signal Processing* (eds. S. Furui and M. M. Sondhi), Marcel Dekker, New York, pp. 701-737 (1991).
 - [50] A. E. Rosenberg, "The Use of Cohort Normalized Scores for Speaker Verification," Proc. Int. Conf. Spoken Language Processing, Banff, Th.sAM.4.2, pp. 599-602 (1992).
 - [51] M. Savic and S. K. Gupta, "Variable Parameter Speaker Verification System Based on Hidden Markov Modeling," Proc. IEEE Int. Conf. Acoust., Speech, Signal Processing, S5.7, pp. 281-284 (1990).
 - [52] A. Setlur and T. Jacobs, "Results of a Speaker Verification Service Trial Using HMM Models," EUROSPEECH'95, Madrid, pp. 639-642 (1995)
 - [53] K. Shikano, "Text-Independent Speaker Recognition Experiments Using Codebooks in Vector Quantization," J. Acoust. Soc. Am. (abstract), Suppl. 1, No. 77, p. S11 (1985).
 - [54] M. -H. Siu, G. Yu and H. Gish, "An Unsupervised, Sequential Learning Algorithm for the Segmentation of Speech Waveforms with Multiple Speakers," Proc. IEEE Int. Conf. Acoust. Speech, Signal Processing, San Francisco, pp. I-189-192 (1992).
 - [55] F. K. Soong, A. E. Rosenberg, and B. -H. Juang, "A Vector Quantization Approach to Speaker Recognition," AT&T Technical Journal, No. 66, pp. 14-26 (1987).
 - [56] F. K. Soong and A. E. Rosenberg, "On the Use of Instantaneous and Transitional Spectral Information in Speaker Recognition," IEEE Trans. Acoust. Speech, Signal Processing, Vol. ASSP-36, No. 6, pp. 871-879 (1988).
 - [57] M. Sugiyama, "Segment Based Text Independent Speaker Recognition," Proc. Spring Meeting of Acoust. Soc. Japan (in Japanese), pp. 75-76 (1988).
 - [58] N. Z. Tishby, "On the Application of Mixture AR Hidden Markov Models to Text Independent Speaker Recognition," IEEE Trans. Acoust. Speech, Signal Processing, Vol. ASSP-30, No. 3, pp. 563-570 (1991).
 - [59] L. Wilcox, F. Chen, D. Kimber and V. Balasubramanian, "Segmentation of Speech Using Speaker Identification," Proc. IEEE Int. Conf. Acoust., Speech, Signal Processing, pp. I-161-164 (1994).
 - [60] Y. -C. Zheng and B. -Z. Yuan, "Text-Dependent Speaker Identification Using Circular Hidden Markov Models," Proc. IEEE Int. Conf. Acoust., Speech, Signal Processing, S13.3, pp. 580-582 (1988).