

論文 / 著書情報
Article / Book Information

論題	SPLIT法による特定話者大語彙単語音声認識
著者	菅村 昇, 古井 貞熙
掲載誌/書名	日本電信電話株式会社研究開発本部 研究実用化報告, Vol. 34, No. 12, pp. 1677-1685
発行日	1985,

SPLIT 法による特定話者大語彙 単語音声認識*

管 村 昇
古 井 貞 熙

あらまし 学習用音声データを用いて、スペクトルパターンをベクトル量子化し、数百個の離散化されたスペクトルパターン(擬音韻標準パターンとよぶ)を作成する。この擬音韻標準パターンとこの系列で表現した単語標準パターンを利用し、DP マッチングによって認識を行う新しい単語音声認識法(SPLIT 法)を提案する。この方法は、DP マッチング時のスペクトル距離計算量と単語標準パターンのメモリ量が、単語標準パターンとして分析されたパラメータをそのまま用いる方法に比べ、大幅に少なくて済むという特徴をもち、この特徴は、認識対象単語数が多くなればなるほど有効である。本論文では、この方法について説明するとともに、特定話者大語彙単語音声認識に適用した結果について報告し、提案する方法の有効性を述べる。また、大語彙領域における単語登録のための発声負担を軽減するために、単語標準パターンの自動生成法についても提案する。

1 ま え が き

音声認識の中でも、単語ごとに区切って発声された音声認識する単語音声認識は、特定の使用条件のもとでは実用化されはじめている。単語音声の認識方法としては、DP (ダイナミックプログラミング) マッチングによる方法⁽¹⁾⁽²⁾と、音韻認識による方法⁽³⁾がある。特定話者の単語音声認識の場合、前者では、あらかじめ認識対象単語を発声し、そのまま標準パターンとして登録しておく方法が基本的である。標準パターンとして利用者本人のパターンを用いているため、通常音韻認識による方法より高い認識率が得られる。市販されている単語音声認識装置の大部分では、DP マッチングの手法が用いられている。

しかし、この DP マッチングを用いる認識方法では、入力単語音声の各フレームに対し、単語標準パターンのすべてのフレームとのスペクトル距離を計算する必要があり、この演算量は、ほぼ単語数に比例して増加する。このため大語彙に対しては、能率の良い単語音声認識手法を確立しておく必要がある。本論文では、DP マッチングにおけるスペクトル距離計算量が、認識単語数に依存しない新しい単語音声認識法を提案するとともに、その有効性を実験結果を通して明らかにする。さらに、大語彙単語音声認識における単語標準パターン登録のための発声負担を軽減することを目的として、単語標準パターンの自動生成法についても報告する。

2 SPLIT 法による単語音声認識

DP マッチングを用いる単語音声認識における単語標準パターンの表現法としては、これまで大きく分けて、(1) フレームごとの分析パラメータをそのまま用いる上述の基本的方法⁽¹⁾と、(2) 各単語を音韻の連鎖で表現し、各音韻の継続時間長を設定する方法⁽²⁾がある。これらの方法の概略を図1の(a)、(b)に示す。前者の方法で

* Speaker-Dependent Large Vocabulary Word Recognition Using SPLIT Method. By Noboru SUGAMURA and Sadaoki FURUI.
この研究は情報通信基礎研究部第四研究室で行われたものである。

は、入力単語音声は、各単語標準パターンと直接比較され、入力フレームごとに、すべての単語標準パターンフレームとのスペクトル距離計算が行われる。ここでは、この認識方法を直接マッチング (Direct Matching) と呼ぶことにする。一方、後者では図 1 (b) に示すように、あらかじめ音韻標準パターンを作成しておき、入力フレームは、まず、これらの標準パターンとの間で距離計算が行われ、その後、各単語の音韻継続長に従って DP マッチングが実行される。この認識方法を前者に対比させ、間接マッチング (Indirect Matching) と呼ぶことにする。認識対象語彙数が増加した場合、直接マッチングでは、単語標準パターンのメモリ量、DP マッチング実行時におけるスペクトル距離計算量が、ほぼ単語数に比例して増加する。一方、間接マッチングでは、これらの問題は回避できる長所をもつ反面、複雑に変化するスペクトル時系列を、きわめて少数の音韻標準パターンで表現するため、類似語が多くなる大語彙単語音声認識においては、前者ほどの認識率が期待できないという問題点がある。

本論文では、これら両認識手法の中間的なものとして、ベクトル量子化⁽⁴⁾⁽⁵⁾したスペクトルパターンを用いる新しい認識方法を提案する。この基礎になっているのは、単語辞書の表現法に関して、音韻のように大きな単位を用いるのでは粗すぎるが、分析されたスペクトルパラメータそのものを用いて高精度に表現する必要はないのではないか、という考えかたである。このため、音韻との対応にとられない方法でスペクトルパターンをベクトル量子化し、数百個程度の代表的なスペクトルパターンを作成する。このパターンをここでは、擬音韻標準パターンと呼ぶ。単語標準パターンは、この擬音韻標準パターンの系列で表現される。音韻単位の標準パターンを用いる場合は、音韻の変形を考慮した高度な辞書作成技術が必要であったが、多数のスペクトルパターンを用意することにより、最適な単

語辞書が自動的に作成できるようになる。

このように、単語音声認識にベクトル量子化の手法を持ち込むことにより、間接マッチングの長所 (スペクトル距離計算量とメモリ量が、認識対象語彙数に依存しない) を継承しつつ、かつ高精度な認識が期待できる。

2.1 単語音声認識システム

前述した考え方をもとに、ベクトル量子化手法を用いる新しい単語音声認識法について述べる。単語標準パターンを、ベクトル量子化によって得られる擬音韻標準パターンの系列 (Strings of Phoneme-Like Templates) として、表現することと、擬音韻標準パターンの作成にスペクトルパターン分布空間を分割 (SPLITting) していくことを兼ねて、この認識手法を、SPLIT 法と名付ける⁽⁶⁾。認識システムの構成を図 2 に示し、認識手順について説明する。認識手順は、つぎの 3 つの段階に分けられる。(1) および (2) は、学習段階の処理である。

- (1) 擬音韻標準パターンの作成
- (2) 単語標準パターンの作成
- (3) 単語音声認識

2.1.1 擬音韻標準パターンの作成

未知単語音声入力の認識に先立って、あらかじめいくつかの単語を発声し、これを擬音韻標準パターン作成用のデータとして用いる。いまこの音声データを LPC 分析し、 $P_1, P_2, \dots, P_N$ 個のフレームのパラメータセット (以後パターンと呼ぶ) が得られているとする。それぞれのパターン P は、 n 個の LPC ケプストラムパラメータで構成されているとし、これを p_{ik} ($k=1, 2, \dots, n$) で表す。このような学習データを用いて、擬音韻標準パターンを作成する方法 (ベクトル量子化) について以下に説明する。

[アルゴリズム]⁽⁶⁾

- (1) N 個のパターン P_i のそれぞれに対して P_i とは

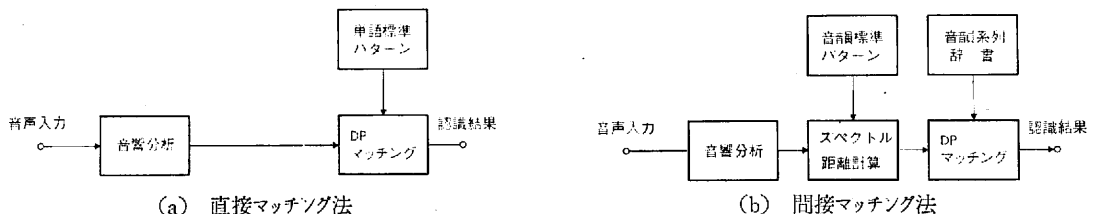


図 1 DP マッチングによる単語音声認識法

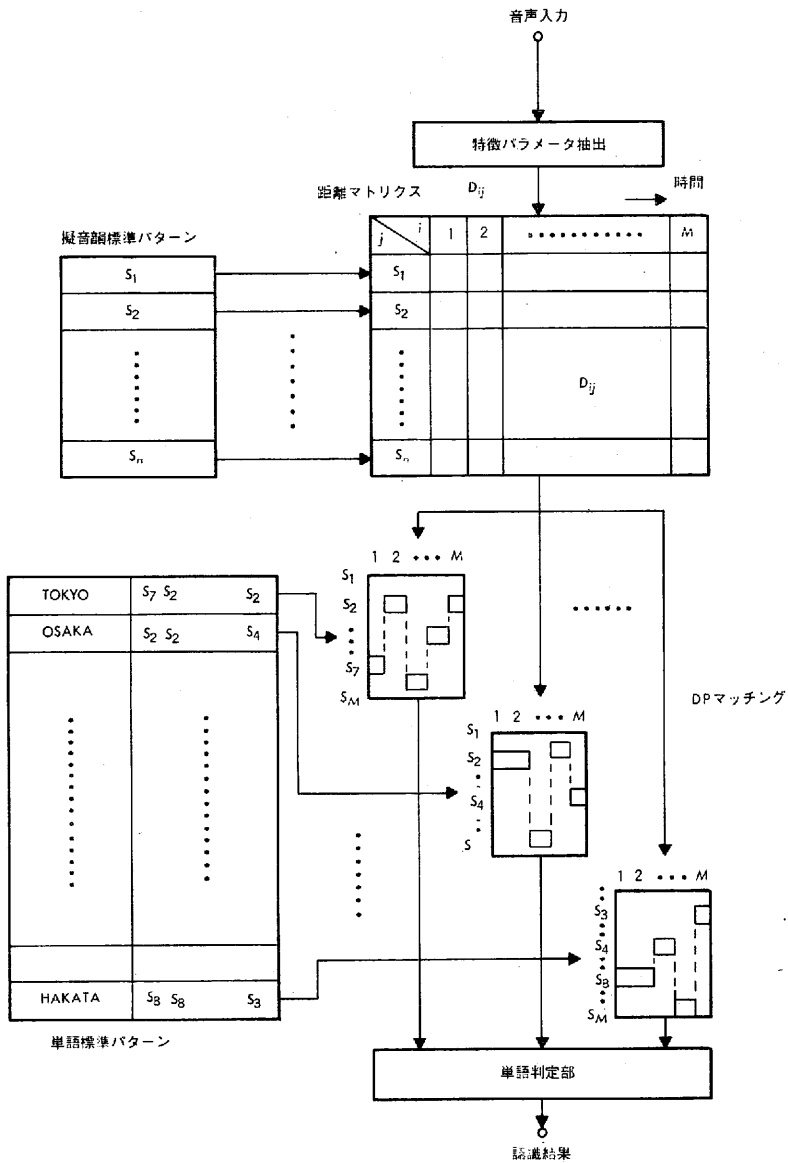


図 2 SPLIT 法による単語音声認識

異なる P_j とのスペクトル距離を次式で計算する。

$$D_{ij} = D(P_i, P_j) \\ = \sum_{k=1}^n (p_{ik} - p_{jk})^2 \quad (i \neq j) \quad (1)$$

(2) $D_{ij} < \theta_s$ を満たすパターン数を各 i について求める。このパターン数を $m_i (i=1, 2, \dots, N)$ とする。こ

こで、 θ_s は、あらかじめ設定しておくスペクトル距離のしきい値であり、 P_i と θ_s 以下のスペクトル距離にあるパターンは、 P_i と類似しているとみなす。

(3) 最大の m_i を有するパターン P_i を見出し、 P_i 自身と P_i から θ_s 以下のスペクトル距離に含まれるすべてのパターンを用いて、これらのパターンの平均値を計算し、擬音韻標準パターン S_i として蓄える。この平均値

の計算は、対数化を行わない線形スペクトル領域で行うため、自己相関係数を用いて平均化を行ったのちケプストラムに変換して蓄える。

(4) (3) で用いたすべてのパターンを、はじめのパターンセットから除去し、(1)–(3) をもとのパターンがなくなるまで繰り返す。

以上のアルゴリズムにおけるパターン間のスペクトル距離計算は、一度だけ実行しておけばよい。2 回目以後は、このスペクトル距離行列を参照し、各パターンごとに θ_s 以下のパターンを検索し、このパターンが、それ以前のどの擬音韻標準パターンにも含まれなかったかどうかだけを調べればよい。このアルゴリズムにおいて与えるべきパラメータは、 θ_s のみで簡単に擬音韻標準パターンを作成することができる。 θ_s を大きく設定することは、粗い量子化を意味しており、 θ_s を小さくすれば、細かいスペクトル表現が可能になる。

2.1.2 単語標準パターンの作成

認識対象単語を発声し、スペクトル分析したパラメータと擬音韻標準パターンとのスペクトル距離計算をフレームごとに行い、スペクトル距離が最小となる擬音韻標準パターンを選択する。このようにして認識対象単語を、擬音韻標準パターンの系列（正確には、擬音韻標準パターンの番号の列）で表現することにより、分析パラメータそのものを用いる直接マッチング法に比べ、単語標準パターン用のメモリ量を大幅に削減することができる。

2.1.3 単語音声認識

あらかじめ準備した擬音韻標準パターンと単語標準パターンを用いて、未知入力単語音声の認識を行う。入力単語音声は、スペクトル分析されフレームごとに擬音韻標準パターンとのスペクトル距離計算が行われ、この結果と単語標準パターンを用いて DP マッチングが実行される。すべての単語標準パターンと DP マッチングが行われた後、マッチング距離が最小となる単語を認識結果として出力する。ここで DP マッチングを行う場合のスペクトル距離計算は、入力単語音声 1 フレームごとに、各擬音韻標準パターンとの距離を 1 回だけ計算しておけば、すべての認識対象単語のマッチングに利用できる。すなわち、SPLIT 法では、DP マッチングに必要なスペクトル距離計算が、認識対象単語数に依存しなくなり、直接マッ

チング法に比べ距離計算量を大幅に低減できる。

2.2 SPLIT 法の特徴

SPLIT 法の特徴を、単語標準パターン用メモリ量、DP マッチング時におけるスペクトル距離計算の観点から、直接マッチング法と比較する。

(1) 単語標準パターン用メモリ量の比較

認識対象単語数を L 、各単語のフレーム数を M_l ($l=1, 2, \dots, L$)、1 フレームのスペクトルパラメータの個数を n 、その精度を N_a ビットとする。また、擬音韻標準パターンの個数を N_s とすると、単語標準パターンおよび擬音韻標準パターンの蓄積に必要なメモリ量は、直接マッチング法では、

$$\sum_{l=1}^L M_l \times n \times N_a \text{ ビット} \quad (2)$$

SPLIT 法では、

$$N_s \times n \times N_a + \left(\sum_{l=1}^L M_l \right) \times (\log_2 N_s) \text{ ビット} \quad (3)$$

となる。

(2) スペクトル距離計算量の比較

入力単語音声 1 フレームあたりのスペクトル距離計算量は、DP マッチングの窓幅を N_w とすると、前述のパラメータを用いて、

直接マッチング法では、

$$N_w \times L \quad (4)$$

SPLIT 法では、

$$N_s \quad (5)$$

となる。

認識対象単語をパラメータにとり、 $N_s=256$ 個、 $N_a=16$ ビット、 $n=16$ パラメータ、 $N_w=15$ フレーム、 $M_l=50$ フレーム（簡単のためにすべての l について一定）とした場合の、距離計算量の比を図 3 に示す。この結果、SPLIT 法は、認識対象単語数が多くなればなるほど、直接マッチング法に比べて有利であることがわかる。

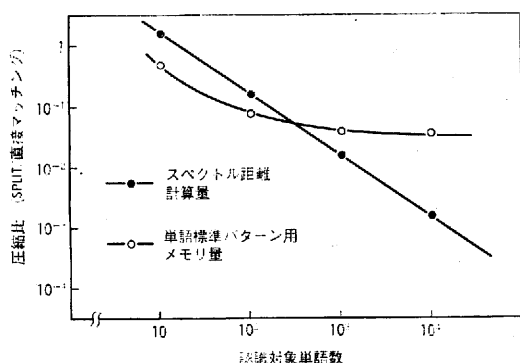


図3 スペクトル距離計算量、単語標準パターン用メモリ量の圧縮比(SPLIT/直接マッチング)と単語数との関係

3 大語彙単語音声認識

前章で説明したように、SPLIT法は、直接マッチング法に比べ、数々の長所を有しているが、認識率が直接マッチング法に比べ大幅に低下するならば有効な方法とはいえない。そこでSPLIT法の有効性が発揮できる大語彙領域において、認識率を評価尺度として直接マッチング法とSPLIT法との比較実験を行った。

3.1 認識実験条件および実験結果

男性4名が、約2週間程度の期間をわいて2回発声した日本の都市名641単語を認識対象として実験を行った。実験の諸条件を表1に示す。

擬音韻標準パターンは、641単語の中からランダムに選んだ64個の単語(約3000フレーム)を用いて作成し、そのうち256パターンを選択し実験に用いた。

1回目の発声データを用いて擬音韻パターンと単語標準パターンを作成し、2回目の発声データを未知入力としたときの認識率と、その逆の場合の認識率とを各話者について求め、4名の話者で平均した結果を図4に示す。横軸は、順位で、縦軸はその順位までに正解が含まれる割合を示している。図には直接マッチング法の場合の結果も合せて示している。スペクトル距離尺度としては、ケプストラム距離(CEPと略記)と重み付き最尤スペクトル距離(WLR⁽⁷⁾)の2種類を用いて行った。この実

表1 実験に用いた音声データおよび分析、認識条件

項 目	内 容
音声データ	国内都市名、641単語 男性4名、各人2回発声 計算機室内(70~85 dBA)で発声 マイクロフォンで収録
分析条件	8 kHz サンプリング、12ビット/サンプル フレーム長 32 ms、ハミング窓 フレーム周期 16 ms
特徴パラメータ	LPC ケプストラムおよび自己相関係数、 各16次
認識条件	両端点フリー DP マッチング スペクトル距離尺度—ケプストラム または WLR

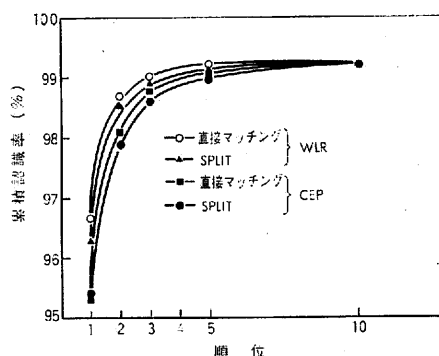


図4 641単語音声認識実験結果の比較(SPLIT法と直接マッチング法の比較およびスペクトル距離尺度の比較)

験結果からつぎのことが明らかとなった。

(1) スペクトル距離尺度として CEP を用いる場合は、SPLIT法と直接マッチング法との第1位の認識率の差はない。

(2) スペクトル距離尺度として WLR を用いた場合、SPLIT法による第1位の認識率は、直接マッチング法のそれに比べて、0.4%低い。これは擬音韻標準パターンを作成するときにスペクトルが平均化され、スペクトルのピーク付近への重み付け効果が小さくなるためと考えられる。

(3) 擬音韻標準パターンの作成用として64個の単語だけを用いたが、これらの擬音韻標準パターンが、それ

以外の単語の認識に対しても有効であることが確認された。

(4) 直接マッチング法において高い認識率が得られる話者は、SPLIT 法によっても高い認識率が得られ、両者の方法による認識率の差は、最大の話者で 1 % であった。SPLIT 法による認識率のみが極端に低下する話者は存在せず、スペクトル情報をベクトル量子化したことによる固有の悪影響は生じなかった。

(5) すべての誤認識単語について、誤認識のタイプをつぎの A-E の 5 つに分類する。A は音声区間の検出エラー、B は母音のエラー (UOZU→O: ZU など)、C は語頭の子音のエラー (SAGA→KAGA)、D は C に含まれない子音のエラー (O: GAKI→O: MACHI など)、E はその他のエラー (MINOO→MINO など) である。この結果、誤認識の中では、音声区間の検出エラーによるものが最も多く、誤認識率 3-4 % のうち 1 % がこのタイプである。これは騒音レベル 70-85 dB (A) の計算機室で収集した音声を用いたためと考えられる。ついで語頭の子音のエラーによるものが多いが、音韻的にきわめて近い単語は、DP マッチングのように、単語全体における平均スペクトルマッチング量で評価する方法では、その認識能力にも限界があり、別の手法⁽⁵⁾を併用するなどの工夫が必要である。

3.2 擬音韻標準パターン数と認識率との関係

前項の実験では 256 個の擬音韻標準パターンを用いたが、この数をさらに減少させた場合の認識率の変化を実験的に検討した。具体的な擬音韻標準パターン数としては、16, 32, 64, 128 の場合を選択した。章 2 で説明したベクトル量子化アルゴリズムにおいて、 θ_c を 8 種類設定し、擬音韻標準パターンセットを作成した。この 8 種類の擬音韻標準パターンセットを用いて、それぞれから所定のパターン数 (たとえば 16 個) を選択し、単語標準パターンを作成した。このときの全学習単語の全フレームのスペクトル誤差の総和を計算し、その値が最も小さくなるときの擬音韻標準パターンセットを 8 種類のセットから選び、認識実験に用いた。

男性話者 4 名について、各パターン数 (16, 32, 64,

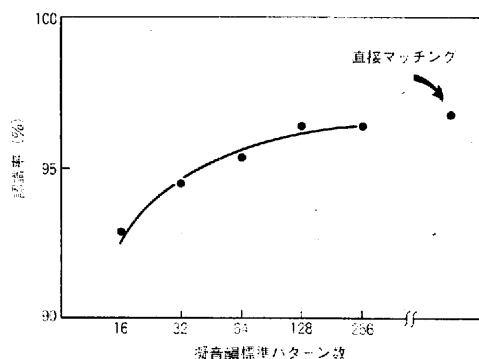


図 5 擬音韻標準パターン数と認識率との関係

128) における擬音韻パターンを選択し、2 回目の発声データを未知入力とした認識実験を行った。4 名の平均認識率を図 5 に示す。この結果、認識率の低下は、擬音韻標準パターンの減少に対して緩やかであり、64 個にしても 95 % 以上の認識率が得られることが明らかになった。さらに、日本語における音韻の数よりも少ない 16 個にしても 93 % 程度の認識率が確保できるという興味深い結果を得た。

4 単語標準パターンの自動生成

前章までに特定話者、大語彙単語音声認識の一手法として SPLIT 法を提案し、認識実験を通してその有効性を確認してきた。すなわち、SPLIT 法を用いれば、認識率を直接マッチング法に比べほとんど低下させることなく、単語標準パターンのメモリ量、DP のためのスペクトル距離計算量を大幅に低減できる。しかし、SPLIT 法においても、単語標準パターンの作成には、認識対象単語の全発声を前提としており、大語彙領域では、発声の負担が大きい。そこで、教師付き学習を前提として、擬音韻標準パターン系列で表現された CV 音節パターンの結合により、大語彙単語音声認識における初期単語標準パターンの自動生成について検討した結果について述べる。初期単語標準パターンの「初期」の意味は、最初の認識では、この標準パターンを用いるが、この標準パターンを常時利用するというのではなく、最終的には実際に発声され、正しく認識された音声から作成した標準パ

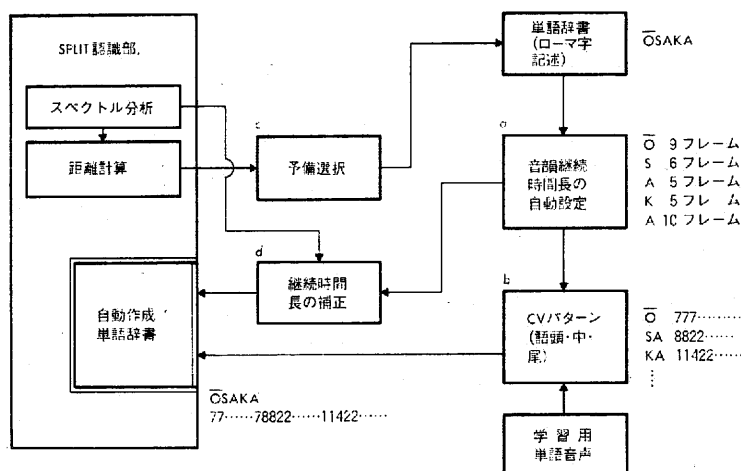


図6 SPLIT法をベースとした自動作成単語標準パターンによる単語音声認識システム
(数字は擬音韻標準パターン番号)

ターンに置き換えるということである。

4.1 CV音節パターンの結合による単語標準パターンの自動生成

単語標準パターン作成のための利用者の発声負担を軽減するために、音韻や音節標準パターンなどをもとにして、単語標準パターンを自動生成する研究が試みられている⁽⁹⁾⁽¹⁰⁾。ここでは、擬音韻標準パターンの系列で表現されたCV音節標準パターンの結合による単語標準パターンの自動生成法を提案する。図2に示したSPLIT法に、この自動生成を付加した認識システム全体の構成を図6に示し、以下にその方法について説明する。(1)、(2)が学習段階の処理であり、(3)、(4)が認識段階の処理である。

(1) 音韻継続長の設定(図6-a)

CV音節パターンの結合には、規則合成の研究で提案された音韻継続時間長の制御規則を⁽¹¹⁾適用する。この規則は、音韻固有基本長をもとに、モーラ数、前後の音韻による影響などを考慮したものである。この規則を用いることにより、ローマ字記述された各単語内の各音韻の継続時間長が自動的に決定される。

(2) CV音節パターンの作成(図6-b)

認識に先立ち、学習用音声データを用いて、擬音韻標準パターン系列で表現されたCV音節パターンの作成

を行う。CV音節パターンとしては、各CV音節について、語頭、語中、語尾の3種類のパターンを単語中から切りだし蓄積しておく。

(3) 予備選択(図6-c)

擬音韻標準パターンの中には、音韻パターンとの対応が明確なものが存在している。これらのパターンを利用して、入力単語音声に含まれる母音系列を推定し、自動生成すべき単語候補を絞る。このようにすることにより自動生成のための演算量の削減を図るとともに、CV音節パターンのあいまいさに伴う誤認識を減少させることができる。

(4) 音韻継続時間長の補正(図6-d)

(1)における規則で生成された各音韻継続時間長の総和は、同一入力単語の時間長とは一般に一致しない。このためDPマッチングの窓制限からはみ出し、正しくマッチングできない場合が生じる。そこで、入力単語長と各音韻継続時間長の総和との差を、各音韻継続長に、母音部2、子音部1の割合で割り振り、各音韻継続長の補正を行う。

4.2 本手法の特徴

ここで提案した手法の長所としては、(1)単語中から切りだした音節パターンを使用することにより、スペクトル変化情報を保存している。(2)SPLIT法をベースにし

表 2 自動作成辞書による認識実験結果

音声データ	単語数 100 単語(平均)	641 単語
1 回目発声	98.0%	91.9%
2 回目発声	96.7%	89.2%

ているため、DP マッチング実行時のスペクトル距離演算量が少ない。(3) 単語標準パターンの表現法は、ローマ字記述でよく、このため語彙の追加が容易であり、メモリ量も少なくて済む。一方、短所としては、(1) 入力単語音声長への適合が必要であり、発声が終わるまで処理を開始できない。(2) CV 音節パターンを作成するための発声とその処理過程が必要である。

4.3 認識実験方法および実験結果

章 3 の実験で用いた男性 4 名のうち 1 名についてのみ実験を行った。1 回目発声データの約 1/5 を用いて、128 個の擬音韻標準パターンの作成、641 単語中に依存する CV 音節パターンの切りだし、擬音韻標準パターンと母音との対応付けなどを行った。スペクトル距離尺度としては WLR、DP には Staggered Array DP (DP 3-5)⁽¹²⁾ を用いた。認識対象としては、1, 2 回目発声の 641 単語と、これを 6 分割した 100 単語 6 セットを用いた。

認識実験結果を表 2 に示す。641 単語の場合の認識率は、本人の発声した音声を用いる場合に比べ 6% 程度低下するが、単語中の CV 音節パターンを用いれば、90% 程度の認識率は、確保できる見通しが得られた。誤認識単語の多くは、子音部の誤りに起因するものであり CV 音節パターンのバリエーションの記述、結合法などを改良する必要がある。

5 む す び

本論文では、SPLIT 法という新しい単語音声認識方法について、その基本的な考え方、特徴を従来の代表的な認識方法と比較して述べ、大語彙単語音声認識実験を通じて、認識性能を明らかにした。この結果、擬音韻標準パターン数を 64-256 個程度に選べば、SPLIT

法は、直接マッチング法に比べ、DP のためのスペクトル距離計算量、単語標準パターン用のメモリ量が大幅に少なくて済み、しかも、認識精度はほとんど差がないことが明らかになった。本方法は、特定話者大語彙単語音声認識に極めて有効な方法であると言える。本論文では、さらに大語彙単語音声認識における、利用者の単語登録の負担を軽減することを目的として、単語標準パターンの自動生成法を提案し、その有効性を明らかにした。

謝 辞

本研究を進めるにあたって熱心に御討論頂いた名古屋大学板倉文忠教授(元情報通信基礎研究部第四研究室長)に深謝します。音声データの収集に御尽力された音声入出力方式研究室箱田主任研究員に感謝します。また、認識実験に関している御協力、御助言頂いた情報通信基礎研究部第四研究室鹿野主幹研究員、杉山研究主任に感謝します。単語標準パターンの自動生成については、同句坂主任研究員のプログラムを使用させていただきました。ここに記して、厚く御礼申し上げます。

文 献

- (1) 迫江・千葉：動的計画法を利用した時間正規化に基づく連続音声認識、音学誌、27, No. 9, pp. 483~490, 1971.
- (2) 好田・橋本・斎藤：数字音声の機械認識系、信学論 (D), 55-D, No. 3, pp. 186~193, 1972.
- (3) 牧野・本間・城戸：音素を単位とした単語音声認識、音声研資, S 83-18, pp. 134~141, 1983.
- (4) Y. Linde, A. Buzo and R. M. Gray: An Algorithm for Vector Quantizer Design, IEEE Trans., COM-28, No. 1, pp. 84~95, 1980.
- (5) 管村・板倉：スペクトルパターンマッチングによる音声情報圧縮、信学論 (A), 65-A, No. 8, pp. 834~841, 1982.
- (6) 管村・古井：擬音韻標準パターンによる大語彙単語音声認識、信学論 (D), 65-D, No. 8, pp. 1041~1048, 1982.
- (7) 杉山・鹿野：ピークに重みをおいた LPC スペクトルマッチング尺度の提案、音声研資, S 80-13, pp. 101~108, 1980.

- (8) 松永・鹿野・相川・杉山：Top-down 処理と Bottom up 処理を融合した音声認識，音声研資，S 83-49, pp. 381~388, 1983.
- (9) A. E. Rosenberg, L. R. Rabiner, J. G. Wilpon and D. Kahn, Demisyllable-Based Isolated Word Recognition System, IEEE Trans., ASSP **31**, No. 3, pp. 713~726, 1983.
- (10) 箱田和雄：音韻を単位とする単語音声認識における辞書の自動作成法，音響学会講論集，pp. 395~396, 1980.
- (11) 匂坂・東倉：音韻固有の性質を考慮した音韻継続時間長設定，信学研究会，EA 80-66, pp. 51~58, 1981.
- (12) 鹿野・相川：Staggered Array DP マッチング，音声研資，S 82-15, pp. 113~120, 1982.

(1985, 5, 1 受付)

管村 昇



複合通信研究所主任研究員
昭和51年入社。主に音声分析合成，音声認識の研究に従事。現在，連続音声認識の研究に従事。
昭和49年大阪大学工学部電気工学科卒業。51年同大学院電気工学専攻修士課程修了。60年工学博士（大阪大学）。
電子通信学会・日本音響学会・情報処理学会・AVIRG 会員。
電子通信学会56年度業績賞，論文賞，および59年度日本音響学会佐藤論文賞受賞。

古井 貞熙



基礎研究所主幹研究員
昭和45年入社。主に音声認識，話者認識の研究に従事。現在，音声知覚に関する基礎研究に従事。
昭和43年東京大学工学部計数工学科卒業。45年同大学院修士課程修了。53年工学博士（東京大学）。
電子通信学会・日本音響学会・IEEE 会員。
昭和49年度電子通信学会米澤賞および59年度日本音響学会佐藤論文賞受賞。