

(19) 日本国特許庁(JP)

(12) 特 許 公 報(B2)

(11) 特許番号

特許第3698511号  
(P3698511)

(45) 発行日 平成17年9月21日(2005.9.21)

(24) 登録日 平成17年7月15日(2005.7.15)

(51) Int.Cl.<sup>7</sup>

F I

G 1 O L 15/14

G 1 O L 3/00 5 3 5 A

G 1 O L 15/10

G 1 O L 3/00 5 3 1 P

G 1 O L 15/20

G 1 O L 3/00 5 6 1 A

G 1 O L 15/28

G 1 O L 3/00 5 3 1 G

請求項の数 10 (全 18 頁)

|           |                        |           |                     |
|-----------|------------------------|-----------|---------------------|
| (21) 出願番号 | 特願平8-330595            | (73) 特許権者 | 000000295           |
| (22) 出願日  | 平成8年12月11日(1996.12.11) |           | 沖電気工業株式会社           |
| (65) 公開番号 | 特開平10-171489           |           | 東京都港区虎ノ門1丁目7番12号    |
| (43) 公開日  | 平成10年6月26日(1998.6.26)  | (74) 代理人  | 100090620           |
| 審査請求日     | 平成15年1月6日(2003.1.6)    |           | 弁理士 工藤 宣幸           |
|           |                        | (72) 発明者  | 伊達 正晃               |
|           |                        |           | 東京都港区虎ノ門1丁目7番12号 沖電 |
|           |                        |           | 気工業株式会社内            |
|           |                        | 審査官       | 石丸 昌平               |

最終頁に続く

(54) 【発明の名称】 音声認識方法

(57) 【特許請求の範囲】

## 【請求項1】

入力された音声データを認識して処理する音声認識方法において、

音響モデルを構成するHMMの任意の状態間の遷移のしやすさを表す状態遷移制約情報を作成する参照テーブル作成手段と、入力音声データに対する認識結果候補と共に認識ゆう度及び局所ゆう度を算出する認識処理手段と、入力音声データの棄却判定に用いる参照ゆう度を算出する参照ゆう度算出手段とを備え、

上記参照テーブル作成手段で、音響モデルを構成するHMMの状態に対するクラスタリングを行い、生成された状態クラスタにおける各状態間の遷移接続に基づいて状態クラスタ間の遷移確率を算出し、HMMの各状態がどの状態クラスタに属するかを示すヘッダ情報を付加して状態遷移制約情報とすると共に、

10

上記認識処理手段で算出した局所ゆう度及び上記参照テーブル作成手段で作成した状態遷移制約情報を用いて参照ゆう度算出手段により参照ゆう度を算出し、この参照ゆう度と上記認識処理手段で算出した認識ゆう度との比較により、入力音声データの棄却判定を行うことを特徴とする音声認識方法。

## 【請求項2】

請求項1に記載の音声認識方法において、

上記参照ゆう度算出手段で、入力音声データの各フレーム毎に、上記局所ゆう度と上記状態遷移制約情報との加重和を最大化するHMMの状態とその最大値を算出し、このときの最大値を局所参照ゆう度として、各局所参照ゆう度を全フレームについて累積加算する

20

ことによって、上記参照ゆう度を算出することを特徴とする音声認識方法。

【請求項 3】

入力された音声データを認識して処理する音声認識方法において、

音響モデルを構成する HMM の任意の状態間の遷移のしやすさを表す状態遷移制約情報を作成する参照テーブル作成手段と、入力音声データに対する認識結果と共に局所ゆう度及び部分仮説累積ゆう度を算出する認識処理手段と、入力音声データ中の不要語または未知語を処理するために用いる不要語仮説累積ゆう度を算出する不要語処理手段とを備え、

上記認識処理手段で算出した局所ゆう度及び部分仮説累積ゆう度と、上記参照テーブル作成手段で作成した状態遷移制約情報とを用いて、不要語処理手段により不要語仮説累積ゆう度を算出し、この不要語仮説累積ゆう度を用いて入力音声データ中の不要語または未知語の部分を検出し、それ以外の部分の認識を行うことを特徴とする音声認識方法。

10

【請求項 4】

請求項 3 に記載の音声認識方法において、

上記不要語処理手段で、

入力音声データの各フレーム毎に、上記局所ゆう度と上記状態遷移制約情報との加重和を最大化する HMM の状態とその最大値を算出して、このときの最大値を局所参照ゆう度とし、

上記部分仮説累積ゆう度に対するフレーム番号の次フレーム番号を始端フレーム番号とし、この始端フレーム番号にあらかじめ設定した複数個の異なる固定フレーム数を加算して当該個数の終端フレーム番号を得て、

20

上記局所参照ゆう度に補正定数を加算した値を、始端フレーム番号から終端フレーム番号まで累積加算することによって、複数個の異なる不要語仮説ゆう度を算出し、

この不要語仮説ゆう度に上記部分仮説累積ゆう度を加算することによって、上記部分仮説累積ゆう度のフレーム番号に対して、上記複数個の異なるフレーム番号における不要語仮説累積ゆう度を算出することを特徴とする音声認識方法。

【請求項 5】

請求項 3 に記載の音声認識方法において、

上記認識処理手段で行う、不要語仮説に後続する発話内容の部分仮説に対するゆう度計算を、上記不要語仮説累積ゆう度を初期値とし、その不要語仮説累積ゆう度に対するフレーム番号の次フレーム番号を始端フレーム番号として行うことを特徴とする音声認識方法

30

【請求項 6】

請求項 3 に記載の音声認識方法において、

上記参照テーブル作成手段で、音響モデルを構成する HMM の状態に対するクラスタリングを行い、生成された状態クラスタにおける各状態間の遷移接続に基づいて状態クラスタ間の遷移確率を算出し、HMM の各状態がどの状態クラスタに属するかを示すヘッダ情報を付加して状態遷移制約情報とすることを特徴とする音声認識方法。

【請求項 7】

請求項 1 又は 6 に記載の音声認識方法において、

上記状態クラスタ間の遷移確率を、任意の状態クラスタに属する状態から他の状態クラスタに属する状態への遷移接続に対する状態遷移確率に基づいて定義することを特徴とする音声認識方法。

40

【請求項 8】

請求項 1 又は 6 に記載の音声認識方法において、

任意の状態クラスタ  $i$  に属する状態から他の状態クラスタ  $j$  に属する状態への遷移接続に対する状態遷移確率の総和  $f_{i,j}$  ( $i \neq j$ ) と、

上記任意の状態クラスタ  $i$  の内部における遷移接続に対する状態遷移確率の総和を、当該任意の状態クラスタ  $i$  から他のすべての状態クラスタへの遷移接続の束の数で割った値  $f_{i,i}$  を用いて、

他のすべての状態クラスタ  $j$  についての  $f_{i,j}$  ( $i = j$  の場合を含む) の総和に対する

50

、ある  $f_{ij}$  ( $i = j$  の場合を含む) の割合により、状態クラス  $i$  から他の状態クラス  $j$  への遷移確率  $P_{ij}$  を定義することを特徴とする音声認識方法。

【請求項 9】

請求項 1 又は 6 に記載の音声認識方法において、

上記状態クラス  $i$  間の遷移確率を、任意の状態クラス  $i$  に属する状態から他の状態クラス  $j$  に属する状態への遷移接続の個数に基づいて、定義することを特徴とする音声認識方法。

【請求項 10】

請求項 1 又は 6 に記載の音声認識方法において、

上記状態クラス  $i$  間の遷移確率を、上記任意の状態クラス  $i$  に属する状態から他の状態クラス  $j$  に属する状態への遷移接続の有無に基づいて、定義することを特徴とする音声認識方法。

【発明の詳細な説明】

【0001】

【発明の属する技術分野】

本発明は音声認識装置に用いる音声認識方法に関するものである。

【0002】

【従来の技術】

文献名 (1) : 波辺 隆夫、塚田 聡 “音節認識を用いたゆう度補正による未知発話のリジェクション” 電子情報通信学会論文誌 D-II Vol. J75-D-II No.12 1992 年12月 PP. 2002-2009

文献名 (2) : 大河内 正明 “Hidden Markov Model に基づいた音声認識” 日本音響学会誌 42 巻 12 号 (1986)

音声認識装置では、高い認識精度とリアルタイム処理を実現するため、装置が受理できる単語や文法規則等をあらかじめ規定することによって、認識対象を制約して認識処理を行う。しかし、利用者が実際に装置を使用する場合には、認識対象外の発話、言い誤り、言い直し等は避けられない。そこで、ある発話に対する認識結果の信頼性が低い場合には、発話を棄却するリジェクト機能が必要になる。

【0003】

リジェクト機能を付加するための方法として、従来、上記文献 (1) に開示される方法がある。この方法では、音声を表現するモデル (一般に、音響モデルと呼ばれる) として、音素や音節などのサブワード単位の HMM (Hidden Markov Model : 隠れマルコフモデル) を用いることを前提としている。HMM を用いた音声認識方法の詳細については、上記文献 (2) に開示されている。サブワードモデルを連結することによって、認識対象として規定された単語や文などの発話内容の仮説に対するモデルを構成し、各仮説に対するモデルが入力音声データを生成する確率 (ゆう度) を計算する。この計算により、最大ゆう度を与えるモデルに対応する仮説を認識結果とする。

【0004】

この認識手法にリジェクト機能を付加するためには、以上のような認識対象を制約して行うゆう度計算 (認識処理) の他に、入力音声を任意の音素列あるいは音節列として認識するためのゆう度計算を行う。それぞれのゆう度計算の結果得られた最大ゆう度の差を求め、しきい値判定により入力発話のリジェクト判定を行う。

【0005】

【発明が解決しようとする課題】

しかしながら、以上述べた従来方法では、以下に述べる問題がある。

【0006】

(a) トライフォンモデル等のコンテキスト依存音素モデルは、音素コンテキストに依存した異音を表現でき、比較的高い認識精度が得られるため、音響モデルとしてよく用いられる。しかし、音響モデルとして、このトライフォンモデル等のコンテキスト依存音素 (あるいは音節) モデルを用いる場合、リジェクト機能を付加すると、処理量が大幅に増

10

20

30

40

50

加する。このため、リジェクト機能を付加することが困難であった。また、自然な発話に対する音声認識装置の頑健性を向上させるには、入力発話中の不要語や未知語に対処する必要があるが、音声認識に対する処理量と精度において十分な性能を得ることは困難であった。

【0007】

(b) 音響モデルとして、音素や音節などのサブワード単位のモデルを用いない場合（例えば、単語や文節などの単位を用いる場合）、前述した従来の方法は適用できない。

【0008】

(c) 入力発話の一部に不要語や未知語を含む場合、認識のための処理量の増加を抑えた状態で、不要語や未知語部分を検出し、それ以外の発話部分を精度よく認識することは困難である。上記(a)と同様に、音響モデルとしてコンテキスト依存音素（あるいは音節）モデルを用いる場合は、特に処理量が大幅に増加する。

10

【0009】

【課題を解決するための手段】

上記課題を解決するために、本発明に係る音声認識方法は、入力された音声データを認識して処理する音声認識方法において、音響モデルを構成するHMMの任意の状態間の遷移のしやすさを表す状態遷移制約情報を作成する参照テーブル作成手段と、入力音声データに対する認識結果候補と共に認識ゆう度及び局所ゆう度を算出する認識処理手段と、入力音声データの棄却判定に用いる参照ゆう度を算出する参照ゆう度算出手段とを備え、上記参照テーブル作成手段で、音響モデルを構成するHMMの状態に対するクラスタリングを行い、生成された状態クラスタにおける各状態間の遷移接続に基づいて状態クラスタ間の遷移確率を算出し、HMMの各状態がどの状態クラスタに属するかを示すヘッダ情報を付加して状態遷移制約情報とすると共に、上記認識処理手段で算出した局所ゆう度及び上記参照テーブル作成手段で作成した状態遷移制約情報を用いて参照ゆう度算出手段により参照ゆう度を算出し、この参照ゆう度と上記認識処理手段で算出した認識ゆう度との比較により、入力音声データの棄却判定を行うことを特徴とする。

20

【0010】

以上のように、局所ゆう度と状態遷移制約情報を用いて参照ゆう度を算出することで、参照ゆう度の算出に要する演算が加算と大小比較だけになり、棄却判定機能の付加による処理量の増加を小さくすることができる。また、参照ゆう度と認識ゆう度との比較により、  
入力音声データの棄却判定を行うため、認識のための処理量をほとんど増加させることなく、効率的に棄却判定を行うことが可能になる。

30

【0011】

また、他の発明に係る音声認識方法は、入力された音声データを認識して処理する音声認識方法において、音響モデルを構成するHMMの任意の状態間の遷移のしやすさを表す状態遷移制約情報を作成する参照テーブル作成手段と、入力音声データに対する認識結果と共に局所ゆう度及び部分仮説累積ゆう度を算出する認識処理手段と、入力音声データ中の不要語または未知語を処理するために用いる不要語仮説累積ゆう度を算出する不要語処理手段とを備え、上記認識処理手段で算出した局所ゆう度及び部分仮説累積ゆう度と、上記参照テーブル作成手段で作成した状態遷移制約情報とを用いて、不要語処理手段により不要語仮説累積ゆう度を算出し、この不要語仮説累積ゆう度を用いて入力音声データ中の不要語または未知語の部分を検出し、それ以外の部分の認識を行うことを特徴とする。

40

【0012】

これにより、不要語仮説ゆう度の算出に要する演算は加算と大小比較だけになり、不要語処理の付加による処理量の増加を最小限に抑えることができる。

【0013】

【発明の実施の形態】

次に本発明の実施形態を添付図面に基づいて説明する。

【0014】

[第1の実施形態]

50

## 〔構成及び機能〕

本実施形態に係る音声認識方法に用いる音声認識装置の基本構成を図１に示す。

## 【００１５】

図１中の１０は音声分析部である。この音声分析部１０は、入力音声データＤ１０を音響特徴パラメータ時系列Ｄ１１に変換する。具体的には、ＬＰＣ（Linear Predictive Coding）分析等の分析手法を用いて、入力音声データＤ１０を数ｍｓ～数十ｍｓ程度の短時間周期（以後「フレーム」と呼ぶ）毎の音響特徴パラメータに変換する。ここで音響特徴パラメータとは、音声データのスペクトル包絡情報を表現するパラメータで、例えば、ケプストラム（対数スペクトルを逆フーリ工変換した量）やその時間変化量などである。フレーム単位に得られる音響特徴パラメータを音響特徴パラメータ時系列Ｄ１１とする。音声分析部１０で変換された音響特徴パラメータ時系列Ｄ１１は、認識処理部１３に入力される。なお、上記入力音声データＤ１０は、マイクロフォンなどから入力された音声（アナログ信号）をデジタル信号に変換した信号である。

10

## 【００１６】

１１は音響モデルである。この音響モデル１１は、音声表現するモデル（ＨＭＭ）の集合である。モデルの言語的な単位としては、音声の任意の構成要素（音素、音節、単語、文節など）を採用することが可能である。また、音素や音節などのサブワードを単位として採用した場合、コンテキスト独立／依存のどちらのモデルでも使用することができる。つまり、リジェクト機能を付加するために使用する音響モデルが制限されることはない。本実施形態では、例としてトライフォンモデルを使用する場合について説明する。トライ

20

## 【００１７】

１２は言語モデルである。この言語モデル１２は、音声認識装置が受理可能な単語や文法規則（構文）等を規定して、認識対象を制約するモデルである。例えば図２に示すように、有限状態オートマトンを用いて受理可能な単語系列を構文ネットワークの形で記述したものである。

## 【００１８】

１３は認識処理部である。この認識処理部１３では、音声が入力される以前に、即ち認識処理を開始する以前に、音響モデル１１及び言語モデル１２を用いて受理可能な発話内容の仮説を表現するＨＭＭネットワークを構成しておく。このＨＭＭネットワークとは、単語の音素表記や文法規則等の制約に従ってトライフォンモデルを連結して作成する、文字通りＨＭＭのネットワークである。これは、例えば図２に示したような構文ネットワークにおいて、単語の部分をトライフォンモデルの連結によって作成した単語モデル（ＨＭＭ）に置き換えたものである。このようなネットワークを構成することによって、認識処理を効率化することができる。各発話内容の仮説に対応するモデルは、ＨＭＭネットワークの一部として表現される。音声認識装置に発話が入力されると、ＨＭＭネットワークを用いて各仮説に対応するモデルが音響特徴パラメータ時系列Ｄ１１を生成する確率（ゆう度）を計算する。ＨＭＭネットワーク中で最大ゆう度を与える仮説を探索し、認識結果候補Ｄ１２とする。また、このときの最大ゆう度を対数化した最大対数ゆう度を、認識ゆう度Ｄ１３とする。ここで、各仮説に対するゆう度計算は、音響特徴パラメータ時系列Ｄ１１のフレームに同期して並列に行う。各フレームではＨＭＭの各状態に対する出力確率分布計算（当該フレームの音響特徴パラメータを出力する確率の計算）を行い、これを対数化して局所ゆう度Ｄ１４とする。認識ゆう度Ｄ１３は局所ゆう度Ｄ１４とＨＭＭの状態遷移確率を用いて、前述した従来技術文献（２）に開示されるViterbi アルゴリズム等の手段により算出する。

30

40

## 【００１９】

この認識処理部１３で算出された認識結果候補Ｄ１２と認識ゆう度Ｄ１３はリジェクト判定部１６に出力される。局所ゆう度Ｄ１４は参照ゆう度算出部１４に出力される。

## 【００２０】

50

参照ゆう度算出部 14 では、認識処理部 13 で算出される局所ゆう度 D14 と、参照テーブル 15 に格納されている状態遷移制約情報 D15 を用いて参照ゆう度 D16 を算出する。

【0021】

参照テーブル 15 は、参照ゆう度算出部 14 で用いる状態遷移制約情報 D15 を格納しているテーブルである。状態遷移制約情報 D15 は、あらかじめ音響モデル 11 を用いて作成しておく。その作成方法は後述する。

【0022】

リジェクト判定部 16 では、認識ゆう度 D13 と参照ゆう度 D16 を用いて認識結果候補 D12 に対してリジェクト判定を行い、認識結果 D17 を出力する。

【0023】

[音声認識方法]

次に、上記構成の音声認識装置を用いた音声認識方法を説明する。

【0024】

[第1段階]

まず、音響モデル 11 を構成するすべてのトライフォンモデルを用いて、HMM の状態に対するクラスタリングを行う。クラスタリングにより生成される各クラスタを、以後「状態クラスタ」と呼ぶ。クラスタリングにおける距離尺度は各状態を表現するパラメータを用いて定義する。例えば、各状態の出力確率分布が多次元正規分布で表されている場合、多次元正規分布の平均ベクトル（あるいは分散ベクトル）を用いて以下のように定義する。

【0025】

2つの平均ベクトル（あるいは、さらに分散ベクトルを付加したベクトル）

$$x = [a_1, a_2, \dots, a_n]$$

$$y = [b_1, b_2, \dots, b_n]$$

に対して、距離尺度 D は

$$D = (a_1 - b_1)^2 + (a_2 - b_2)^2 + \dots + (a_n - b_n)^2$$

となる。

【0026】

クラスタリング方法としては、LBG アルゴリズム等の一般的なクラスタリングアルゴリズムを用いる。ここでは、より簡易な方法を以下に示す。

【0027】

M 個のサンプル集合  $X = \{x_1, x_2, \dots, x_M\}$  をクラスタリングする場合を考える。しきい値  $T_h$  が与えられているものとする。なお、しきい値  $T_h$  の値は実験的に決定される数値である。

【0028】

まず、任意に 1 個のサンプル、例えば  $x_1$  を取り、これをクラスタ中心  $z_1 (= x_1)$  とする。次に  $x_k$  ( $k = 2, 3, \dots, M$ ) を取り、 $z_1$  と  $x_k$  との距離  $D_{1k}$  を計算する。この計算により、 $D_{1k} \leq T_h$  であれば、その  $x_k$  は  $z_1$  を中心とするクラスタに属すると判定する。また、 $D_{1k} > T_h$  であれば、その  $x_k$  を新たなクラスタ中心  $z_2$  とする。同ようにして、残りのサンプル  $x_k$  について  $z_1, z_2$  との距離  $D_{1k}, D_{2k}$  を計算する。そして、この距離  $D_{1k}, D_{2k}$  のいずれかが  $T_h$  より小さければその  $x_k$  はそのクラスタに属するものとし、そうでなければその  $x_k$  を新たなクラスタ中心  $z_3$  とする。

【0029】

以上の操作を、全てのサンプル  $\{x_1, x_2, \dots, x_M\}$  に対して行い、クラスタリングを終了する。

【0030】

[第2段階]

状態クラスタ間の遷移確率を、以下のようにして算出する。

【0031】

まず、状態クラスタ間の遷移確率を定義する。それぞれの状態クラスタに属する各状態

10

20

30

40

50

は、トライフォンモデル上では他の状態に接続されている。例えば、図3に示すように、状態 $S_1$ は状態 $S_2$ に、状態 $S_2$ は状態 $S_3$ にそれぞれ接続されている。また、トライフォンモデルの終端状態 $S_3$ は、次に続き得るトライフォンモデルの始端状態 $S_4$ 、 $S_5$ 、 $S_6$ にそれぞれ接続されている。一般に、あるトライフォンモデルに対して次に続き得るトライフォンモデルは複数存在するので、トライフォンモデルの終端状態は複数の状態に接続されている。さらに、状態の接続関係には、向き（図3で示す矢印）があり、この向きは一方の状態から他方の状態への遷移方向を表している。このときの遷移の起こりやすさとして、状態遷移確率が付与されている。また、各状態には自己ループ遷移を表す接続も存在する。このようなトライフォンモデル上での状態の遷移接続を、状態クラスタに属する各状態に対して適用する。これにより、任意の状態クラスタ間に、構成要素の状態が作る遷移接続の束ができる。図4はこの様子を示した例である。図4において、状態クラスタ1に属する状態 $S_1$ は、状態クラスタ2に属する状態 $S_2$ に接続されており、トライフォンモデルにおいて状態 $S_1$ から状態 $S_2$ への遷移接続（状態遷移確率 $a_{12}$ ）が存在することを意味する。図4では状態クラスタ1に属する状態から、他の状態クラスタに属する状態への遷移接続を示した。なお、一部においては、状態クラスタ1の内部における遷移接続も示した。状態クラスタ間で同一の遷移方向を持つ遷移接続を1つの束としたものが“遷移接続の束”である。この遷移接続の束を用いて、状態クラスタ間の遷移確率を次の(1)～(3)式により定義する。

【0032】

【数1】

$$P_{ij} = \frac{f_{ij}}{\sum_{k=1}^N f_{ik}} \quad \dots(1)$$

$$f_{ij} = \sum_{\substack{u \in i \\ v \in j}} a_{uv} \quad (i \neq j) \quad \dots(2)$$

$$f_{ii} = \frac{\sum_{u \in i} a_{uu} \cdot q_u + r_i}{Z_i} \quad \dots(3)$$

$$i, j = 1, 2, \dots, N$$

$$u, v = 1, 2, \dots, M$$

$i, j$  ( $= 1, 2, \dots, N$ ) : 状態クラスタ番号

$u, v$  ( $= 1, 2, \dots, M$ ) : 状態番号

$P_{ij}$  : 状態クラスタ  $i$  から状態クラスタ  $j$  への遷移確率

$N$  : 状態クラスタの総数

$M$  : 状態の総数

$a_{uv}$  : 状態  $S_u$  から状態  $S_v$  への状態遷移確率

$a_{uu}$  : 状態  $S_u$  の自己ループ遷移確率

$r_i$  : とともに状態クラスタ  $i$  に属する異なる状態間における遷移接続の個数  
(自己ループ遷移接続は対象外)

$z_i$  : 状態クラスタ  $i$  から他の状態クラスタへの“遷移接続の束”の個数

$q_u$  : ある状態クラスタに属する状態  $S_u$  から他の状態クラスタに属する状態への遷移接続の個数

上式において、 $f_{ij}$  ( $i \neq j$ ) は、状態クラスタ  $i$  から状態クラスタ  $j$  への遷移接続の束

10

20

30

40

50

に対する状態遷移確率の総和を表している。ただし、遷移接続が存在しない状態クラスタ間においては、 $f_{ij} = 0$ である。また、 $f_{ii}$ は、状態クラスタ*i*の内部における遷移接続に対する状態遷移確率の総和を、状態クラスタ*i*から他の状態クラスタへの“遷移接続の束”の個数で割った値を表している。

【0033】

以上、説明した式を用いて状態クラスタ間の遷移確率 $P_{ij}$ を算出する。算出した状態クラスタ間の遷移確率 $P_{ij}$ は、対数化して重み係数（定数） $W$ を乗じる。なお、重み係数 $W$ は後述する参照ゆう度算出部14での動作において説明する。

【0034】

以上のようにして得られた $W \cdot \log P_{ij}$ に、トライフォンモデルの各状態がどの状態クラスタに属するかを示すヘッダ情報を付加して、状態遷移制約情報D15とする。

10

【0035】

<上記以外の定義方法>

なお、状態クラスタ間の遷移確率は、上述した定義方法以外に、以下に示す定義方法によっても可能である。いずれの方法においても、状態クラスタ間の遷移接続に基づいて定義している点が共通している。

【0036】

(a) 状態クラスタ*i*から状態クラスタ*j*への遷移接続の束（遷移接続の束を構成する遷移接続の個数は、1個でも構わない）が存在するか否かによって $f_{ij}$ を定義する。

【0037】

20

これは上記(1)～(3)式において、 $f_{ij}$ 及び $f_{ii}$ を次のように定義し直すことによって得られる。

【0038】

1) 状態クラスタ*i*から状態クラスタ*j*への遷移接続の束が存在するならば、 $f_{ij} = 1$ とする。

【0039】

2) 状態クラスタ*i*から状態クラスタ*j*への遷移接続の束が存在しないならば、 $f_{ij} = 0$ とする。

【0040】

3)  $f_{ii} = 1$ とする。

30

【0041】

(b) 状態クラスタ*i*から状態クラスタ*j*への遷移接続の束を構成する遷移接続の個数により $f_{ij}$ を定義する。

【0042】

これは上記(1)～(3)式において、 $a_{uv}$ 及び $a_{uu}$ を

$$a_{uv} = a_{uu} = 1$$

と定義し直すことによって得られる。

【0043】

[参照ゆう度算出部14の動作]

参照ゆう度算出部14では、次の(4)、(5)式を用いて参照ゆう度D16を算出する。

40

【0044】

【数2】

$$L_G = \sum_{t=1}^T L_g(t) \quad \dots(4)$$

$$L_g(t) = \max_{v \in V} \{W \cdot \log c_{uv} + \log b_v(x_t)\} \quad \dots(5)$$

$$t=1, 2, \dots, T$$

ここで、

$t = 1$  のとき、 $c_{uv} = 1$

$t \neq 1$  のとき、 $c_{uv} = P_{ij}$

ただし、 $u = i, v = j$

また

$c_{uv} = 0$  のとき、 $\log c_{uv} = \text{INH}$

とする。

【0045】

$t$  : フレーム番号

$L_g$  : 参照ゆう度D16

$L_g$  : 対数ゆう度

$T$  : フレーム総数

$W$  : 状態遷移制約情報に対する重み係数

$u$  : フレーム番号 ( $t - 1$ ) において、(5) 式の右辺の最大値を与える状態番号

$v$  : 任意の状態番号

$P_{ij}$  : 状態クラス  $i$  から状態クラス  $j$  への遷移確率

$i, j$  : 任意の状態クラス番号

$V_t$  : 認識処理部 13 において、フレーム番号  $t$  に出力確率分布計算を行う状態全体の集合

$b(x_t)$  : 状態  $v$  における音響特徴パラメータ  $x_t$  の出力確率 (密度)

$x_t$  : フレーム番号  $t$  における音響特徴パラメータ

$\text{INH}$  : 状態クラス間の遷移確率を対数化した値の下限值 (定数)

(5) 式において、 $\log b_v(x_t)$  は、認識処理部 13 より局所ゆう度D14 として与えられ、 $W \cdot \log c_{uv}$  は参照テーブル 15 より状態遷移制約情報D15 として与えられる。したがって、参照ゆう度算出部 14 で行う演算は、加算と大小比較のみである。なお、(5) 式において、状態遷移制約情報に対する重み係数  $W$  は、 $\log c_{uv}$  と  $\log b_v(x_t)$  の  $L_g(t)$  に寄与する割合を調節するためのパラメータ (定数) であり、定数  $\text{INH}$  は状態クラス間の遷移確率を対数化した値の下限值を設定するためのパラメータ (定数) である。ともにその値は実験的に決定する。

【0046】

なお、参照ゆう度  $L_g$  は、任意の発話内容を表現するモデルに対する累積対数ゆう度を表す。また、 $L_g(t)$  は、任意の発話内容を表現するモデルに対する各フレームにおける局所的な対数ゆう度を表す。

【0047】

(5) 式において、フレーム番号 ( $t - 1$ ) における (5) 式の右辺の最大値を与える状態番号を  $u$  とする。 $\log c_{uv}$  は状態番号  $u$  が何であるかによって、次フレーム番号  $t$  における (5) 式の右辺の最大値を与える状態番号の候補を制約する。状態番号  $u$  の状態から状態番号  $v$  の状態への遷移の起こりやすさを制約として用いている。このような状態遷移制約によって、トライフォンモデルが有する音声の時間構造を考慮した参照ゆう度D16の算出を可能にしている。

【0048】

[リジェクト判定部 16 の動作]

リジェクト判定部 16 では、次の (6) 式により入力音声データD10 のリジェクト判定を行う。(6) 式において、リジェクト判定のしきい値  $\theta$  は実験的に決定される。しきい値  $\theta$  の値によって、入力認識対象である場合の認識率と、認識対象外である場合のリジェクト率が変化する。一般に、両者はトレードオフの関係にあるので、所望の性能に合わせしきい値  $\theta$  の値を決定する。

【0049】

$L_M = (L_g - L_R) / T \quad \dots \dots (6)$

$L_R$  : 認識ゆう度D13

10

20

30

40

50

$L_G$  : 参照ゆう度D16

T : フレーム総数

この式により得られた値  $L_M$  としきい値  $\theta$  との大きさを比較してリジェクト判定を行う。  
なお、 $\theta$  はリジェクト判定のしきい値である。

【 0 0 5 0 】

$L_M > \theta$  ならば、入力のリジェクトするように判定し、認識結果D17 として、入力のリジェクトされたことを示す情報を出力する。

【 0 0 5 1 】

$L_M \leq \theta$  ならば、入力のリジェクトせずに、認識結果D17 として、認識結果候補D12 を出力する。

10

【 0 0 5 2 】

[ 効果 ]

以上のように、入力発話のリジェクト判定に用いる参照ゆう度D16 を、認識ゆう度D13 の算出過程で得られる局所ゆう度D14 と、あらかじめ作成した状態遷移制約情報D15 に基づいて算出する。これにより、参照ゆう度D16 の算出に要する演算は加算と大小比較だけになり、リジェクト機能の付加による処理量の増加をきわめて小さくすることができる。また、上記の参照ゆう度D16 は、音響モデル 1 1 が有する音声の時間構造を考慮しつつ、種々の音響的事象に対処可能な定式化を行って算出しているため、音素あるいは音節認識を用いる従来の方法と同等のリジェクト精度を得ることができる。この結果、認識対象外の発話（認識対象語以外の語、あるいは文法外の発話）や、せき、くしゃみなどの非言語音、あるいはベルなどの環境音が装置に入力された場合に、認識のための処理量をほとんど増加させることなく、効率的にリジェクト判定を行うことが可能になる。

20

【 0 0 5 3 】

換言すると、次のようになる。

【 0 0 5 4 】

( a ) 音響モデルとして、音素や音節などのサブワードに対するコンテキスト依存モデルを用いても、リジェクト機能の付加による処理量の増加はほとんどなく、処理機能を高い状態に維持することができる。

【 0 0 5 5 】

( b ) 音響モデルとして、いかなる言語的単位（音素、音節、単語、文節など）のモデルを用いても、リジェクト機能を付加することが可能である。

30

【 0 0 5 6 】

( c ) 音素あるいは音節認識を用いる方法（従来法）と同等のリジェクト精度を得ることができる。

【 0 0 5 7 】

[ 第 2 の実施形態 ]

上記第 1 の実施形態では、認識対象語以外の語や文法外の入力発話を、全体として棄却する方法について説明したが、本実施形態では、入力発話のうち不要な部分だけを部分的に棄却する方法について説明する。即ち、本実施形態では、入力発話の一部に「あのー」、「えーと」等に代表される間投詞や、「○○かな」、「○○とか」等の不要な語、あるいは「（じょ）情報」といったような言いよどみなどを含む場合に対処する方法について説明する。以後、間投詞、不要な語、言いよどみ等を不要語と呼ぶ。

40

【 0 0 5 8 】

本実施形態に係る音声認識装置の基本構成を図 5 に示す。図 5 において入力音声データD20 は、マイクロフォンなどから入力された音声（アナログ信号）をデジタル信号に変換した信号である。入力音声データD20 は、音声分析部 2 0 で音響特徴パラメータ時系列D21 に変換され、認識処理部 2 3 に入力される。

【 0 0 5 9 】

認識処理部 2 3 では音響モデル 2 1、言語モデル 2 2 及び不要語処理部 2 4 を用いて認識処理を行い、入力音声データD20 に対する認識結果D26 を出力する。

50

## 【 0 0 6 0 】

また、不要語処理部 2 4 では、認識処理部 2 3 で算出される局所ゆう度 D22 と部分仮説累積ゆう度 D23、さらに参照テーブル 2 5 に格納されている状態遷移制約情報 D24 を用いて、不要語仮説累積ゆう度 D25 を算出する。不要語仮説累積ゆう度 D25 は認識離処理部 2 3 に出力され、認識結果 D26 の算出に用いられる。

## 【 0 0 6 1 】

以下、本実施形態に係る音声認識方法に特有の機能を中心に説明し、第 1 実施形態と同様の部分については、その説明を省略する。

## 【 0 0 6 2 】

## [ 言語モデル 2 2 ]

言語モデル 2 2 は、第 1 実施形態同様に、音声認識装置が受理可能な単語や文法規則（構文）等を規定して、認識対象を制約するモデルである。例えば、図 6 に示すように有限状態オートマトンを用いて、受理可能な単語系列を構文ネットワークの形で記述したものである。ただし、本実施形態の言語モデル 2 2 では、入力発話中の不要語に対処するため、構文ネットワークの各ノードに、自己遷移として、不要語を表現するアークを付加している。このようにすることによって、任意の単語間において、不要語を受理することが可能になる。

## 【 0 0 6 3 】

## [ 認識処理部 2 3 ]

認識処理部 2 3 の機能は、上記第 1 の実施形態の認識処理部 1 3 とほぼ同様である。ただし、本実施形態の認識処理部 2 3 では、不要語の部分については HMM による明示的な不要語モデルを用意せず、後述する不要語処理部 2 4 を介して各ノードに自己遷移させるようになっている。つまり、不要語処理部 2 4 を不要語モデルとして用いる。このようなネットワークを構成することによって、認識処理及び不要語処理を効率的に行う。各々の発話内容の仮説に対応するモデルは、仮説の任意の単語間において、不要語を受理可能な形で、HMM ネットワークの一部として表現される。音声認識装置に発話が入力されると、HMM ネットワークを用いて、各仮説に対応するモデルが音響特徴パラメータ時系列 D21 を生成する確率（ゆう度）を計算する。各 HMM ネットワーク中で最大ゆう度を与える仮説を探索し、認識結果 D26 とする。

## 【 0 0 6 4 】

入力発話の一部に不要語が含まれる場合の認識結果 D26 は次のようになる。

## 【 0 0 6 5 】

## [ 図 6 の言語モデルを用いた場合 ]

入力発話：「それじゃあー 東京の（こ）交通情報」

認識結果 D26：「# 東京 # 交通情報」（# は不要語を表す記号）

認識結果 D26 に対応する最大ゆう度を対数化した最大対数ゆう度を、以後、「認識ゆう度」と呼ぶ。ここで、各仮説に対するゆう度計算は、音響特徴パラメータ時系列 D21 のフレームに同期して並列に行う。各フレームでは HMM の各状態に対する出力確率分布計算（当該フレームの音響特徴パラメータを出力する確率の計算）を行い、これを対数化して局所ゆう度 D22 とする。

## 【 0 0 6 6 】

認識ゆう度は、局所ゆう度 D22 と HMM の状態遷移確率を用いて、前述した従来技術文献（2）に開示されている Viterbi アルゴリズム等の手段により算出する。ただし、不要語の部分は、不要語処理部 2 4 によりゆう度計算を行う。

## 【 0 0 6 7 】

不要語処理部 2 4 におけるゆう度計算方法は後述する。認識ゆう度を算出する上での不要語処理部 2 4 の扱いは、他の単語モデルと同様である。音響特徴パラメータ時系列 D21 のフレーム番号 1 から任意のフレーム番号までの“発話内容の部分仮説”に対する累積対数ゆう度を、部分仮説累積ゆう度 D23 とする。部分仮説累積ゆう度 D23 は、その部分仮説の終端フレーム番号を付加して不要語処理部 2 4 に出力される。

10

20

30

40

50

## 【 0 0 6 8 】

## [ 不要語処理部 2 4 ]

不要語処理部 2 4 では、次の ( 7 ) ~ ( 9 ) 式を用いて不要語仮説累積ゆう度 D25 を算出する。以下では、不要語を表す発話内容の部分仮説を「不要語仮説」と、この不要語仮説に対する対数ゆう度を「不要語仮説ゆう度」と呼ぶ。不要語仮説ゆう度の算出における始端フレーム番号及び終端フレーム番号をそれぞれ  $t_0$ 、 $t_1$  とし、このときの不要語仮説ゆう度を  $L_g(t_1)$  で表す。不要語仮説累積ゆう度 D25 は、フレーム番号 ( $t_0 - 1$ ) における部分仮説累積ゆう度 D23 と、不要語仮説ゆう度  $L_g(t_1)$  との和として定義する。したがって、不要語仮説累積ゆう度 D25 はフレーム番号 1 からフレーム番号  $t_1$  ( 不要語仮説ゆう度の算出における終端フレーム番号 ) までの発話内容の部分仮説に対する累積対数ゆう度を表している。

10

## 【 0 0 6 9 】

次に、不要語仮説ゆう度  $L_g(t_1)$  について説明する。( 8 ) 式は任意の始端フレーム番号  $t_0$  に対して、異なる ( $T_{del} + 1$ ) 個の終端フレーム番号  $t_1$  ( $= t_0 + T_{min}$ ,  $t_0 + T_{min} + 1$ , ...,  $t_0 + T_{min} + T_{del}$ ) における不要語仮説ゆう度  $L_g(t_1)$  を算出することを表している。

## 【 0 0 7 0 】

なお、( 8 ) 式において  $L_g(t_1)$  は、各フレームにおける不要語仮説に対する局所的な対数ゆう度を表す。また、補正定数  $R$  は、不要語仮説ゆう度の変域を調節するためのパラメータ(定数)で、その値は実験的に決定される。不要語仮説ゆう度  $L_g(t_1)$  は、 $[L_g(t) + R]$  を、始端フレーム番号  $t_0$  から終端フレーム番号  $t_1$  まで、累積加算したものと定義する。不要語仮説ゆう度  $L_g(t_1)$  の算出に用いる値で、(最小フレーム数 - 1)  $T_{min}$  と、最大フレーム数及び最小フレーム数の差  $T_{del}$  の値は、実験的に決定する。

20

## 【 0 0 7 1 】

また、 $L_g(t)$  を定義する ( 9 ) 式において、 $\log b_v(x_t)$  は、認識処理部 2 3 より局所ゆう度 D22 として与えられ、 $W \cdot \log c_{uv}$  は、参照テーブル 2 5 より状態遷移制約情報 D24 として与えられる。ここで、状態遷移制約情報に対する重み係数  $W$  は、 $\log c_{uv}$  と  $\log b_v(x_t)$  の  $L_g(t)$  に寄与する割合を調節するためのパラメータ(定数)であり、定数  $INH$  は状態クラスタ間の遷移確率を対数化した値の下限值を設定するためのパラメータ(定数)である。ともにその値は実験的に決定する。

30

## 【 0 0 7 2 】

これにより、不要語処理部 2 4 で行う演算は、加算と大小比較のみとなる。

## 【 0 0 7 3 】

次に、( 9 ) 式における  $\log c_{uv}$  の働きについて説明する。

## 【 0 0 7 4 】

フレーム番号 ( $t - 1$ ) において、( 9 ) 式の右辺に最大値を与える状態番号を  $u$  とする。 $\log c_{uv}$  は状態番号  $u$  が何であるかによって、次フレーム番号  $t$  において ( 9 ) 式の右辺の最大値を与える状態番号の候補を制約する。状態番号  $u$  の状態から状態番号  $v$  の状態への遷移の起こりやすさを制約として用いている。このような状態遷移制約によって、ライフォンモデルが有する音声の時間構造を考慮した不要語仮説ゆう度の算出を可能にしている。不要語処理部 2 4 で算出されたフレーム番号  $t_1$  における不要語仮説累積ゆう度 D25 は、認識処理部 2 3 に出力される。認識処理部 2 3 では、不要語仮説に後続する発話内容の部分仮説に対するゆう度計算を、不要語仮説累積ゆう度 D25 を初期値とし、フレーム番号 ( $t_1 + 1$ ) を始端フレーム番号として行う。このようにすることによって、認識処理部 2 3 における認識ゆう度の計算は、単語仮説(単語を表す発話内容の部分仮説)に対するゆう度と不要語仮説に対するゆう度を同様に扱って行うことができる。

40

## 【 0 0 7 5 】

## 【 数 3 】

$$L_B(t_1) = L_F(t_0 - 1) + L_G(t_1) \quad \dots(7)$$

$$L_G(t_1) = \sum_{t=t_0}^{t_1} \{L_g(t) + R\} \quad \dots(8)$$

$$t_1 = t_0 + T_{min} + y$$

$$y = 0, 1, \dots, T_{del}$$

$$L_g(t) = \max_{v \in V_t} \{W \cdot \log c_{uv} + \log b_v(x_t)\} \quad \dots(9)$$

$$t = 1, 2, \dots, T$$

10

ここで、

$t = 1$  のとき、 $c_{uv} = 1$

$t \neq 1$  のとき、 $c_{uv} = P_{ij}$

ただし、 $u = i, v = j$

また、

$t_0 = 1$  のとき、 $L_F(t_0 - 1) = 0$

$c_{uv} = 0$  のとき、 $\log c_{uv} = I N H$

20

とする。

【 0 0 7 6 】

$t$  : フレーム番号

$L_B(t_1)$  : フレーム番号  $t_1$  における不要語仮説累積ゆう度 D25

$L_F(t_0 - 1)$  : フレーム番号  $(t_0 - 1)$  における部分仮説累積ゆう度 D23

$L_G(t_1)$  : フレーム番号  $t_1$  における不要語仮説ゆう度

$t_0$  : 不要語仮説ゆう度の算出における始端フレーム番号

$t_1$  : 不要語仮説ゆう度の算出における終端フレーム番号

$T_{min}$  : 不要語仮説ゆう度の算出に用いる (最小フレーム数 - 1)

$T_{del}$  : 不要語仮説ゆう度の算出に用いる最大フレーム数と最小フレーム数の差

30

$R$  : 補正定数

$T$  : フレーム総数

$W$  : 状態遷移制約情報に対する重み係数

$u$  : フレーム番号  $(t-1)$  において、(9) 式の右辺に最大値を与える状態番号

$v$  : 任意の状態番号

$i, j$  : 任意の状態クラスタ番号

$V_t$  : 認識処理部 2 3 において、フレーム番号  $t$  に出力確率分布計算を行う状態全体の集合

$b_v(x_t)$  : 状態  $v$  における音響特徴パラメータ  $x_t$  の出力確率 (密度)

$x_t$  : フレーム番号  $t$  における音響特徴パラメータ

40

$I N H$  : 状態クラスタ間の遷移確率を対数化した値の下限值 (定数)

【 効果 】

本実施形態では、入力発話の一部に不要語を含む場合に対処するため、不要語仮説ゆう度を、認識ゆう度算出過程で得られる局所ゆう度 D22 と、あらかじめ作成した状態遷移制約情報 D24 に基づいて算出する。これにより、不要語仮説ゆう度の算出に要する演算は加算と大小比較だけになり、不要語処理の付加による処理量の増加を非常に小さくすることができる。この不要語処理の付加による処理量の増加は、HMM による明示的な不要語モデル (一般に、garbage モデルと呼ばれ、種々の不要語を 1 種類の HMM でモデル化するもの) を用いる方法よりも小さく抑えることができる。

【 0 0 7 7 】

50

また、不要語仮説ゆう度は、音響モデル 2 1 が有する音声の時間構造を考慮しつつ、種々の音響的事象に対処可能な定式化を行って算出しているため、高い不要語検出精度を得ることができる。これにより、入力発話の一部に不要語を含む場合に、認識のための処理の増加を低く抑えて効率的に不要語を検出し、不要語以外の部分の認識率を向上させることが可能になる。しかも、不要語が入力発話の先頭、末尾、任意の単語間において複数含まれている場合にも、高い精度で不要語を検出することができる。

【 0 0 7 8 】

換言すると、以下ようになる。

【 0 0 7 9 】

( a ) 音響モデルとして、音素や音節などのサブワードに対するコンテキスト依存モデルを用いても、不要語処理の付加による処理量の増加は非常に小さく、処理機能を高い状態に維持することができる。

10

【 0 0 8 0 】

( b ) 音響モデルとして、いかなる言語的単位 ( 音素、音節、単語、文節など ) のモデルを用いても、処理機能の低下を招くことなく、不要語処理を付加することができる。

【 0 0 8 1 】

( c ) garbage モデルを用いる場合に比べて、不要語を表現するモデルの音響的分解能が高い。したがって、種々の不要語の音響的バリエーションに対処することが可能である。

【 0 0 8 2 】

20

( d ) garbage モデルを用いていないので、不要語を含む多量の音声データを用いてあらかじめ不要語モデルのパラメータ推定 ( 学習 ) をする必要がない。このため、学習用の音声データとして用意しにくい言いよどみ等にも対処することができるようになる。

【 0 0 8 3 】

[ 変形例 ]

( 1 ) 第 2 の実施形態では、入力発話の一部に不要語が含まれる場合に対処する方法について説明したが、本発明の音声認識方法は不要語に限らず、入力発話の一部に未知語 ( 認識対象語以外の語 ) が含まれる場合に対処する方法としても使用可能である。その場合、図 5 における言語モデル 2 2 を、例えば図 7 に示すように記述する。このようにすることによって、入力発話全体を棄却するのではなく、入力発話中の未知語部分を効果的に検出し、未知語以外の部分の認識率を向上させることができる。

30

【 0 0 8 4 】

例えば、図 7 に示した言語モデルを用いて

「ニューヨーク観光情報」

と発声した場合を考える。「ニューヨーク」が未知語である場合、本発明により

「@観光情報」( @ は未知語を表す記号 )

という認識結果を出力することが可能になる。

【 0 0 8 5 】

【 発明の効果 】

以上、詳述したように、本発明の音声認識方法によれば、次のように効果を奏することができる。

40

【 0 0 8 6 】

( 1 ) 局所ゆう度と状態遷移制約情報を用いて参照ゆう度を算出することで、参照ゆう度の算出に要する演算を低減することができ、棄却判定機能の付加による処理量の増加を小さくすることができる。

【 0 0 8 7 】

( 2 ) 参照ゆう度と認識ゆう度との比較により、入力音声データの棄却判定を行うため、認識のための処理量をほとんど増加させることなく、効率的に棄却判定を行うことが可能になる。

【 0 0 8 8 】

50

(3) 不要語仮説ゆう度の算出に要する演算を低減することができ、不要語処理の付加による処理量の増加を最小限に抑えることができる。

【図面の簡単な説明】

【図1】本発明の第1の実施形態に係る音声認識装置の基本構成を示すブロック図である。

【図2】第1の実施形態の言語モデル（構文ネットワーク）例を示す模式図である。

【図3】トライフォンモデルにおける状態の遷移接続の例を示す模式図である。

【図4】状態クラスタ間における遷移接続の例を示す模式図である。

【図5】本発明の第2の実施形態に係る音声認識装置の基本構成を示すブロック図である。

10

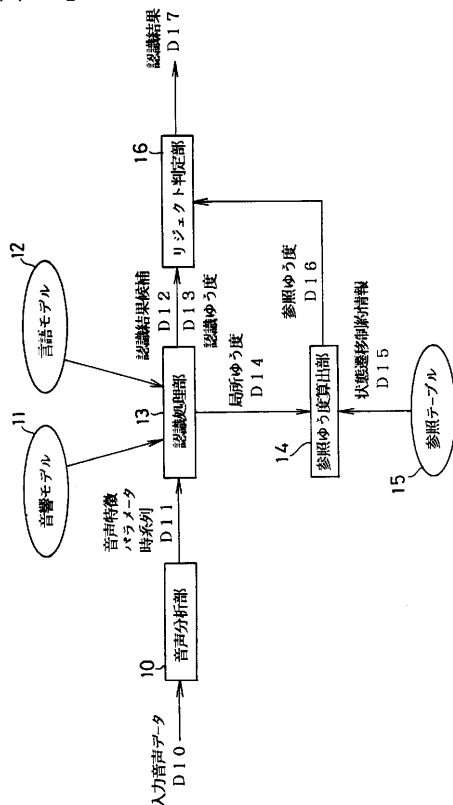
【図6】第2の実施形態の言語モデル（構文ネットワーク）例を示す模式図である。

【図7】本発明の変形例に係る言語モデル（構文ネットワーク）例を示す模式図である。

【符号の説明】

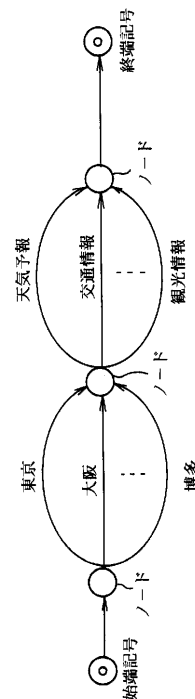
10, 20 : 音声分析部、11, 21 : 音響モデル、12, 22 : 言語モデル、13, 23 : 認識処理部、14 : 参照ゆう度算出部、15, 25 : 参照テーブル、16 : リジェクト判定部、24 : 不要語処理部、D10, D20 : 入力音声データ、D11, D21 : 音響特徴パラメータ時系列、D12 : 認識結果候補、D13 : 認識ゆう度、D14, D22 : 局所ゆう度、D15, D24 : 状態遷移制約情報、D16 : 参照ゆう度、D17, D26 : 認識結果、D23 : 部分仮説累積ゆう度、D25 : 不要語仮説累積ゆう度。

【図1】



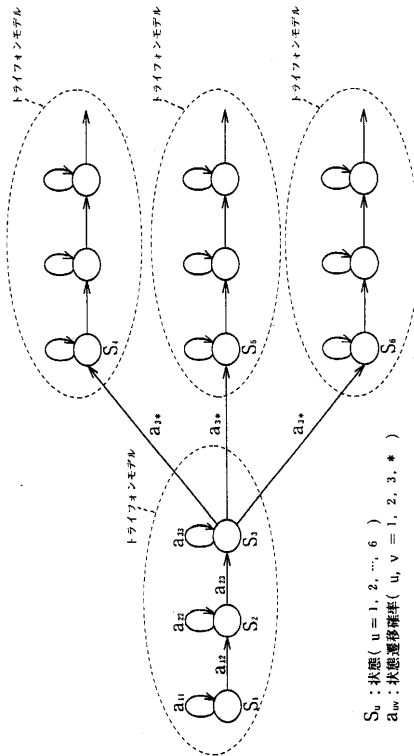
本発明の第1実施形態の基本構成

【図2】



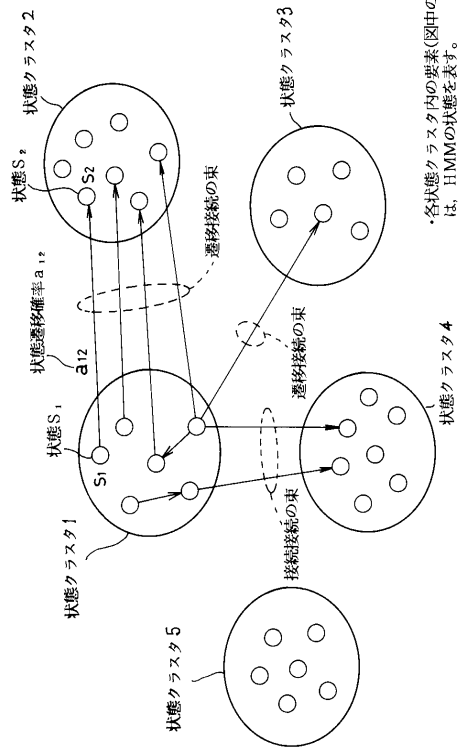
第1の実施形態の言語モデル（構文ネットワーク）例

【 図 3 】



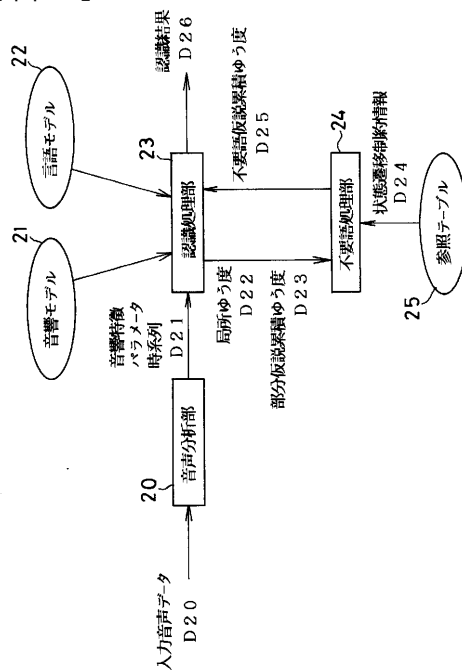
トライフオンモデルにおける状態の遷移接続の例

【 図 4 】



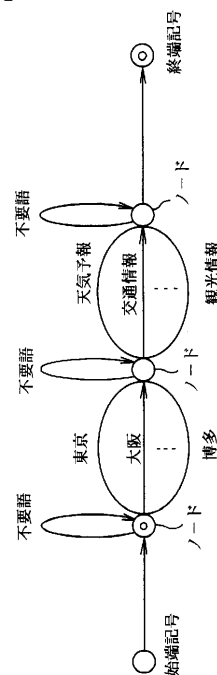
## 状態クラスタ間における遷移接続の例

【 図 5 】



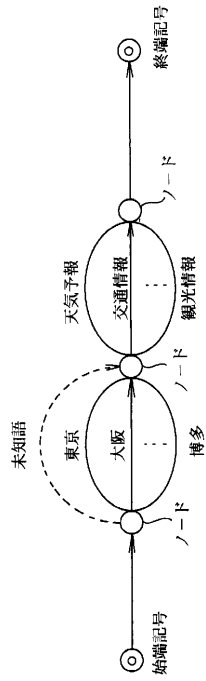
## 本発明の第2実施形態の基本構成

【 図 6 】



第2の実施形態の言語モデル（構文ネットワーク）例

【図 7】



利用形態の説明における言語モデル（構文ネットワーク）例

---

フロントページの続き

(56)参考文献 特開平09-106295(JP,A)

特開平09-062290(JP,A)

特開平09-062292(JP,A)

特開平08-241096(JP,A)

特開平09-081181(JP,A)

特開平10-171489(JP,A)

伊達, 易, 三木, 状態クラスタ間の遷移確率を用いた認識対象外発話のリジェクション法, 日本音響学会平成9年度春季研究発表会講演論文集, 日本, 1997年 3月17日, 2-Q-3, Pages 145-146

(58)調査した分野(Int.Cl.<sup>7</sup>, DB名)

G10L 15/00-15/28

JICSTファイル(JOIS)