

(19) 日本国特許庁(JP)

(12) 特 許 公 報(B2)

(11) 特許番号

特許第4836076号
(P4836076)

(45) 発行日 平成23年12月14日(2011.12.14)

(24) 登録日 平成23年10月7日(2011.10.7)

(51) Int.Cl.

F I

G 1 0 L 15/10 (2006.01)

G 1 0 L 15/10 3 0 0 F

G 1 0 L 15/06 (2006.01)

G 1 0 L 15/06 3 1 0 T

G 1 0 L 15/18 (2006.01)

G 1 0 L 15/06 4 0 0 V

G 1 0 L 15/28 (2006.01)

G 1 0 L 15/18 2 0 0 Z

G 1 0 L 15/28 3 6 0 Z

請求項の数 4 (全 30 頁)

(21) 出願番号 特願2006-47385 (P2006-47385)
 (22) 出願日 平成18年2月23日(2006.2.23)
 (65) 公開番号 特開2007-225931 (P2007-225931A)
 (43) 公開日 平成19年9月6日(2007.9.6)
 審査請求日 平成21年2月20日(2009.2.20)

(出願人による申告)平成17年度独立行政法人情報通信研究機構、研究テーマ「大規模コーパス音声対話翻訳技術の研究開発」に関する委託研究、産業活力再生特別措置法第30条の適用を受ける特許出願

(73) 特許権者 393031586
 株式会社国際電気通信基礎技術研究所
 京都府相楽郡精華町光台二丁目2番地2
 (74) 代理人 100099933
 弁理士 清水 敏
 (72) 発明者 松田 繁樹
 京都府相楽郡精華町光台二丁目2番地2
 株式会社国際電気通信基礎技術研究所内
 (72) 発明者 實廣 貴敏
 京都府相楽郡精華町光台二丁目2番地2
 株式会社国際電気通信基礎技術研究所内
 (72) 発明者 コンスタンティン・マルコフ
 京都府相楽郡精華町光台二丁目2番地2
 株式会社国際電気通信基礎技術研究所内

最終頁に続く

(54) 【発明の名称】 音声認識システム及びコンピュータプログラム

(57) 【特許請求の範囲】

【請求項1】

それぞれ異なる発話環境での発話音声のデコードに最適化された、それぞれ所定の音響特徴量をパラメータとする複数の音響モデル群を記憶するための記憶手段と、

入力される音声から前記所定の音響特徴量を算出するための特徴量算出手段と、

前記特徴量算出手段により算出される前記音響特徴量に基づいて、それぞれ前記複数の音響モデル群の混合重み適応化により、前記入力される音声の発話環境に適応化された複数の適応化音響モデルを作成するためのモデル適応化手段と、

前記複数の適応化音響モデルを用いて、前記入力される音声の前記所定の音響特徴量を音声認識を目的にデコードし音声認識結果の複数の仮説を出力するためのデコード手段と

10

前記デコード手段が出力する前記複数の仮説を、前記複数の仮説内の各単語に対して算出される一般化単語事後確率に基づいて統合し出力するための仮説統合手段とを含み、
 前記仮説統合手段は、

前記デコード手段が出力する前記複数の仮説の各々に対し、各単語の一般化単語事後確率の関数であるスコアを算出するためのスコア算出手段と、

前記複数の仮説から、各単語にスコアが付された単語ラティスを作成するためのラティス作成手段と、

前記単語ラティス内の始点から終点までの経路のうち、当該経路上の単語の各々に対し算出された前記スコアが所定の条件を充足する経路の上の単語列を前記音声認識結果とし

20

て出力するための最適経路探索手段とを含み、

前記仮説の各々の各単語には、入力音声における当該単語の持続時間を特定するための情報が付されており、

前記スコア算出手段は、

前記デコード手段が出力する前記複数の仮説の各々に対し、各単語の一般化単語事後確率を算出するための一般化単語事後確率算出手段と、

前記一般化単語事後確率算出手段により算出された一般化単語事後確率と、前記単語ラティス中の各単語の持続時間を特定するための情報との関数として前記スコアを各単語に対し算出するための関数計算手段とを含む、音声認識システム。

【請求項 2】

10

前記最適経路探索手段は、前記単語ラティス内の始点から終点までの経路のうち、当該経路上の単語の各々に対し算出された前記スコアの和が最大となる経路の上の単語列を前記音声認識結果として出力するための最大スコア経路探索手段を含む、請求項 1 に記載の音声認識システム。

【請求項 3】

前記関数計算手段は、以下の式により前記スコアを各単語に対して算出するための手段を含む、請求項 1 又は請求項 2 に記載の音声認識システム。

【数 1】

$$A_n = T_n \log C_n$$

20

ただし、 A_n はある仮説中の n 番目の単語の前記スコア、 T_n は入力音声における当該単語の持続時間、 C_n は当該単語に対して前記一般化単語事後確率算出手段により算出された一般化単語事後確率を、それぞれ示す。

【請求項 4】

コンピュータにより実行されると、当該コンピュータを、請求項 1 ~ 請求項 3 のいずれかに記載の音声認識システムとして動作させる、コンピュータプログラム。

【発明の詳細な説明】

【技術分野】

【0001】

この発明は大語彙の連続音声認識装置及び方法に関し、特に、雑音に強く、発話スタイルの変動に対しても頑健に音声を認識することが可能な連続音声認識システムに関する。

30

【背景技術】

【0002】

近年、雑音又は発話スタイルに対して頑健な音声認識の研究が盛んに行なわれている。実環境において音声認識を使用するためには、通行する自動車等の乗り物から発せられるエンジン雑音や風切り音、駅、オフィス内等の人の声、コンピュータからのファンの音等、多種多様な雑音環境において高精度な音声認識が実現されなければならない。

【0003】

さらに雑音だけでなく、使用者の年齢や性別、また感情や体調によってその発話スタイルは刻一刻と変化する。音声認識装置は、そのような発話スタイルの変動に対しても雑音と同様に頑健でなければならない。

40

【0004】

雑音又は発話スタイル等個別の変動に対する頑健化手法が従来から数多く提案されてきた。これについては後掲の非特許文献 1 を参照されたい。本明細書では以下、音声の音響的言語的特徴に影響する要因のことを総じて「発話環境」と呼ぶこととする。

【0005】

雑音に対して頑健な音響特徴量の分析手法として、「SS (Spectrum Subtraction) 法」(後掲の非特許文献 2 を参照されたい。)を音声認識の前処理として用いる手法が提案されている。これ以外にも、RASTA (Relative Spectra)、DMFCC (Differential Mel Frequency

50

Cepstrum Coefficient)等、いくつかの音響分析手法が提案されている。

【0006】

SS法では、雑音重畳音声のスペクトルに対して雑音スペクトルを減算することにより、SNR(信号対雑音比)を改善している。RASTA法では、個々の周波数バンドの値の変化に対して、音声情報が多く含まれている1から12Hzの変調スペクトラム成分を抽出することにより雑音の影響を軽減している。またDMFCCはFFT(高速フーリエ変換)によって得られるフーリエ係数に対して、隣り合う係数間で差分をとり、音声等のピッチを持つスペクトルを強調することによって耐雑音性を改善している。

【0007】

雑音に頑健な音響モデルの研究としては、PMC(Parallel Model Combination)法(後掲の非特許文献5を参照されたい。)、ヤコビ適応法(後掲の非特許文献6を参照されたい。)、MLLR(Maximum Likelihood Linear Regression)(後掲の非特許文献7を参照されたい。)による雑音適応等が提案されている。

【0008】

これらのうち、PMC法は、HMM(隠れマルコフモデル)の出力確率分布を線形スペクトル領域に変換し雑音スペクトルを重畳することにより、環境雑音への適応を行なう手法である。このPMC法につき簡単に説明する。

【0009】

PMC法の概念を図28を参照して説明する。図28を参照して、PMC法の対象となる元の音響モデルが、音響の特徴量からなる音響空間600において領域610の付近に存在する音響をモデル化したものであるものとする。このとき、音声認識対象の雑音を含んだ音声データ領域612は、雑音のために元の領域610からずれたものとなる。そこで、領域612と領域610との差分を考え、この差分に相当する量を音響モデル610に加えることにより音響モデルの音響空間600内における位置を領域612まで移動するよう音響モデルを変換する。

【0010】

このようにして変換した後の音響モデルを用いれば、領域612の付近に存在する雑音を含んだ音声については、元の音響モデルを用いたものより高い精度で認識できる。

【0011】

ヤコビ適応法は、雑音の変化に伴う出力確率分布の非線形変換を線形近似することにより、雑音環境へ高速に適応する手法である。

【0012】

MLLRを用いた雑音適応は、無雑音音声と雑音重畳音声との間の分布移動を回帰行列を用いて表現し、音響モデル全体を雑音モデルに適応化する手法である。

【0013】

さらに、雑音の分布の時間変動を逐次的に推定することにより、非定常雑音に対する認識精度を改善する手法(後掲の非特許文献9を参照されたい。)が提案されている。

【0014】

発話スタイルに対する頑健性の改善手法としては、発話スタイル依存の音響モデルを用いる手法の他、ロンバード効果によるスペクトルの変形を考慮した手法(非特許文献8を参照されたい。)及び個々の母音HMMの最後に無音状態を追加することにより音声強調発声や言直し発話に頑健な音響モデルを構築する手法(非特許文献10を参照されたい。)等が提案されている。そのほかにも、講演音声等の音素継続時間の短い発声を含む音声に対して、分析フレーム周期又はウィンドウ幅を自動選択することにより認識精度を改善する手法(非特許文献11、12参照)が提案されている。

【0015】

これらの頑健化手法は主として、雑音や発話スタイル等の個別の変動に対する頑健化である。音声認識を実環境で用いるためには、複数の発話環境が刻一刻と変化する状況であ

10

20

30

40

50

っても頑健に音声を認識することができなければならない。このような種々の外乱に対して頑健な音声認識を実現するための方法は大きく二つに分類することができると考えられる。発話環境の変動に頑健な音響モデル及び言語モデルを用いて単数のデコーダで認識を行なうシングルタイプの方法と、お互いに異なる環境に適応化された複数の音響モデル及び言語モデルを使用して得られた複数の仮説を統合するパラレルタイプの手法とである。

【0016】

シングルタイプの音声認識システムを構築するためには、広い発話環境の音声を頑健に認識する音響モデル及び言語モデルが必要である。そのために、男性及び女性双方の学習データから性別独立な音響モデルを推定する等、複数の発話環境のデータを用いてHMMのモデルパラメータ推定を行なうことにより頑健性を改善する手法がある。しかし、男性女性等のお互いの音響的特徴が大きく異なる場合ではなく、種々のSNRのデータを用いて学習する場合、個々の音素モデルの分布が過度に広がることにより音素分類精度の低下が懸念される。従って、このようなモデル化法には頑健化の限界があると考えられる。

【0017】

セグメントモデル（非特許文献13を参照されたい。）では、時間的に離れた音響特徴ベクトル間の相関を計算することで音声の非定常な振舞いのモデル化を試みている。時間的に離れた特徴ベクトル間の相関として発話環境の変動をモデル化することができるならば、セグメントモデルにおいて広い発話環境の音声を頑健に認識できる可能性がある。しかし、効率的な相関の計算方法やモデルパラメータの増大等の問題により十分な精度は得られていない。

【0018】

一方、パラレルタイプによる音声認識は、個々の音響モデルや言語モデルの利用可能な発話環境が限られていたとしても、それらを複数個使用しパラレルにデコーディングすることにより、個々の音素間の分類精度を低下させることなく広い発話環境の音声を頑健に認識できる可能性がある。

【0019】

このような音声認識の例としては、SNRに依存した音響モデルを用いて得られた複数の仮説を最大尤度基準で選択する手法、複数のお互いに異なる音響特徴量を用いて音声認識を行ない、得られた複数の仮説を単語単位で統合する仮説統合法（非特許文献15参照）が提案されている。

【0020】

【特許文献1】特開2005-221678号公報

【特許文献2】特開2005-164837号公報

【非特許文献1】中村、『実音響環境に頑健な音声認識を目指して』、信学技報、EA2002-12、pp.31-36、2002。

【非特許文献2】S.F.ボル、『スペクトル減算を用いた音声の中の音響雑音の抑制』、IEEE音響音声信号処理論文集、第ASSP-27巻、第113-120頁、1979年。（S.F. Boll, 『Suppression of Acoustic Noise in Speech Using Spectral Subtraction,』IEEE Trans. Acoust., Speech, Signal Processing, vol. ASSP-27, pp. 113-120, 1979.）

【非特許文献3】H.ヘルマンスキ及びN.モーガン、「音声のRASTA処理」、IEEE音声及び音響処理トランザクション、第2巻、第4号、第587-589頁（H. Hermansky and N. Morgan, 『RASTA Processing of Speech,』IEEE Trans. Speech and Audio Processing, vol. 2, no. 4, pp. 587-589, 1994.）

【非特許文献4】J.チェン、K.K.パリワル、S.ナカムラ、『頑健な音声認識のための差分パワースペクトル由来のケプストラム』、音声コミュニケーション、第41巻第2-3号、第469-484頁、2003年。（J. Chen, K.K. Paliwal, S. Nakamura, 『Cepstrum Derived from Diffe

10

20

30

40

50

rentiated Power Spectrum for Robust Speech Recognition,」Speech Communication, vol. 41, no. 2-3, pp. 469-484, 2003.)

【非特許文献5】M. ゲールズ及びS. ヤング、『パラレルモデルの組合せを用いた頑健な連続音声認識』、IEEE 音声及び音響処理論文集、第4巻、第5号、第352-359頁、1996年。(M. Gales and S. Young, 『Robust Continuous Speech Recognition Using Parallel Model Combination,』IEEE Trans. on Speech and Audio Processing, vol. 4, No. 5, pp. 352-359, 1996.)

10

【非特許文献6】Y. ヤマグチ、S. タカハシ及びS. サガヤマ、『ヤコビアン適応アルゴリズムを用いた環境雑音への音響モデルの高速適応』、ユーロスピーチ予稿集、97、第2051-2054頁、1997年。(Y. Yamaguchi, S. Takahashi and S. Sagayama, 『Fast Adaptation of Acoustic Models to Environmental Noise Using Jacobian Adaptation Algorithm,』Proc. Eurospeech, 97, pp. 2051-2054, 1997.)

【非特許文献7】C. J. レゲッタ及びP. C. ウッドランド、『連続密度隠れマルコフモデルの話者適応のための最大尤度線形回帰』、コンピュータ音声及び言語、第9巻、第171-185頁、1995年。(C. J. Leggetter and P. C. Woodland, 『Maximum Likelihood Linear Regression for Speaker Adaptation of Continuous Density Hidden Markov Models,』Computer Speech and Language, vol. 9, pp. 171-185, 1995.)

20

【非特許文献8】J. C. ジャンカ、『ロンバード効果とその聴者及び自動音声認識装置に対する役割』、アメリカ音響学会誌、第93巻、第510-524頁、1993年。(J. C. Junqua, 『The Lombard Reflex and its Role on Human Listeners and Automatic Speech Recognizer,』J. Acoustic Soc. Amer., vol. 93, pp. 510-524, 1993.)

30

【非特許文献9】K. ヤオ、B. E. シー、S. ナカムラ及びZ. カオ、『非定常雑音における頑健な音声認識のための連続EMアルゴリズムによる残存雑音の補償』、ICSLP2000予稿集、第1巻、第770-773頁、2000年。(K. Yao, B. E. Shi, S. Nakamura and Z. Cao, 『Residual Noise Compensation by a Sequential EM Algorithm for Robust Speech Recognition in Nonstationary Noise,』Proc. ICSLP2000, vol. 1, pp. 770-773, 2000.)

【非特許文献10】奥田、松井、中村、『誤認識時の言い直し発話における発話スタイルの変動に頑健な音響モデル構築法』信学論、vol. J86-DII, no. 1, pp. 42-51, 2003.

40

【非特許文献11】奥田、河原、中村、『ゆう度基準による分析周期・窓長の自動選択手法を用いた発話速度の補正と音響モデルの構築』信学論、vol. J86-DII, no. 2, pp. 204-211, 2003.

【非特許文献12】南條、河原、『発話速度に依存したデコーディングと音響モデルの適応』信学技報、SP2001-103、2001.

【非特許文献13】M. オステンドルフ、V. ディガラキス及びO. キンバル、『HMMからセグメントモデルへ：音声認識のためのストカスティックモデリングの統一見解』、IEEE 音声及び音響処理論文集、第4巻、第5号、第360-378頁、1996年。

50

(M. Ostendorf, V. Digalakis and O. Kimball, 'From HMMs to Segment Models: A Unified View of Stochastic Modeling for Speech Recognition,' IEEE Trans. Speech and Audio Proc., vol. 4, no. 5, pp. 360 - 378, 1996.)

【非特許文献14】伊田、中村、『雑音GMMの適応化とSN比別マルチパスモデルを用いたHMM合成による高速な雑音環境適応化』信学論、vol. J86-D-II, no. 2, pp. 195 - 203, 2003.

【非特許文献15】K. マルコフ、T. マツイ、R. グルーン、J. ツァン、S. ナカムラ、『DARPA SPINE2用の雑音及びチャネル歪に頑健なASRシステム』、IEICE情報及びシステム論文集、第E86-D巻、第3号、2003年。(K. Markov, T. Matsui, R. Gruhn, J. Zhang, S. Nakamura, 'Noise and Channel Distortion Robust ASR System for DARPA SPINE2 Task,' IEICE Trans. Inf. & Syst., vol. E86-D, no. 3, 2003.)

【非特許文献16】M. オステンドルフ及びH. シンガー、『最大尤度連続状態分割を用いたHMMトポロジー設計』、コンピュータ音声及び言語、第11巻、第1号、第17 - 41頁1997年。(M. Ostendorf and H. Singer, 'HMM Topology Design Using Maximum Likelihood Successive State Splitting,' Computer Speech and Language, vol. 11, no. 1, pp. 17 - 41, 1997.)

【発明の開示】

【発明が解決しようとする課題】

【0021】

しかしながら上述したいずれの方法においても、例えばモデルの変換に時間を要すること、雑音又は発話スタイル等、個別の要素の変動に的確に対応することが難しいこと、等から、実環境における雑音を含んだ音声や、発話スタイルが変動する音声に対して音声認識を精度よく行なうことは未だ可能でない。

【0022】

そこで本件出願人は、特許文献1において、次のような音声認識システムを提案している。すなわち、予め複数種類の音響モデルを準備しておく。入力音声に基づいて、これら音響モデルから、発話環境に適応化された複数の適応化音響モデルをGMM(混合ガウス分布モデル)を用いた高速な適応化処理で作成する。こうして得られた複数の適応化音響モデルを用いて音声認識を行なうことで複数の仮説を作成する。それら仮説を組合わせて一つの単語ラティス(単語グラフ)を作成する。当該ラティスにおいて、音響モデルと言語モデルとから単語ごとに算出される尤度に基づき、経路全体として最も尤度が高くなる経路を探索する。こうして探索された経路は元の仮説を尤度に基づいて統合した仮説となる。複数の仮説に基づいて得た単語ラティスから、最も確からしい経路を組み合わせるので、仮説単独よりも、最終的な結果の精度が高くなる。

【0023】

この音声認識システムによれば、理論の上からも従来の技術と比較してよい仮説が得られるはずであり、実験によってもそうした効果が明らかとなっている。特に、雑音と発話スタイルの変動とに対して頑健な音声認識が得られている。

【0024】

しかし、この特許文献1に開示の技術でも、音声認識精度にはさらに改善の余地がある可能性がある。特に音声認識は種々の自然言語処理の入り口となるため、音声認識精度をできるだけ高めることが重要である。

【0025】

それゆえにこの発明の目的は、雑音等の個別の変動に実時間で追従して、従来技術より

10

20

30

40

50

もさらに精度高く認識することができる音声認識システムを提供することである。

【 0 0 2 6 】

この発明の他の目的は、雑音等の個別の変動だけでなく、発話スタイルの変動に対しても頑健に、従来技術よりもさらに精度高く音声認識することができる音声認識システムを提供することである。

【課題を解決するための手段】

【 0 0 2 7 】

本発明の第 1 の局面に係る音声認識システムは、それぞれ異なる発話環境での発話音声のデコードに最適化された、それぞれ所定の音響特徴量をパラメータとする複数の音響モデル群を記憶するための記憶手段と、入力される音声から所定の音響特徴量を算出するための特徴量算出手段と、特徴量算出手段により算出される音響特徴量に基づいて、それぞれ複数の音響モデル群の混合重み適応化により、入力される音声の発話環境に適応化された複数の適応化音響モデルを作成するためのモデル適応化手段と、複数の適応化音響モデルを用いて、入力される音声の所定の音響特徴量を音声認識を目的にデコードし音声認識結果の複数の仮説を出力するためのデコード手段と、デコード手段が出力する複数の仮説を、各単語に対して算出される一般化単語事後確率に基づいて統合し出力するための仮説統合手段とを含む。

【 0 0 2 8 】

複数の音響モデル群に対し、モデル適応化手段による混合重み適応化が適用され、入力される音声の発話環境に適応化された適応化音響モデル群が作成される。これら適応化音響モデル群を用いて入力される音声の音響特徴量をデコード手段がデコードし、複数の仮説が出力される。これら複数の仮説が互いに相補的である場合、すなわち、ある仮説の誤った部分が他の仮説では正しく音声認識されている場合、統合手段により仮説を統合することにより、より精度の高い音声認識結果を得られる可能性が高い。仮説統合手段において、これら仮説を統合する際の基準として、各単語に対し算出される一般化単語事後確率を使用すると、一般化単語事後確率を使用しない場合よりも音声認識精度が高くなることが確認できた。

【 0 0 2 9 】

好ましくは、仮説統合手段は、デコード手段が出力する複数の仮説の各々に対し、各単語の一般化単語事後確率の関数であるスコアを算出するためのスコア算出手段と、複数の仮説から、各単語にスコアが付された単語ラティスを作成するためのラティス作成手段と、単語ラティス内の始点から終点までの経路のうち、当該経路上の単語の各々に対し算出されたスコアが所定の条件を充足する経路の上の単語列を音声認識結果として出力するための最適経路探索手段とを含む。

【 0 0 3 0 】

複数の仮説の各々に対し、各単語のスコアを一般化単語事後確率の関数としてスコア算出手段で算出する。これらスコアが付与された仮説から、単語ラティスを作成する。この単語ラティスの始点から終点までの経路のうちで、当該経路上の単語の各々に対し算出されたスコアが、経路全体として所定の条件を充足するものを探索する。こうして、複数の仮説から一般化単語事後確率に基づくスコア付きの単語ラティスを作成し、スコアが所定の条件を充足する経路上の単語列を最終的な音声認識結果とすることで、一般化単語事後確率を使用しない場合よりも音声認識精度が高くなることが確認できた。

【 0 0 3 1 】

より好ましくは、仮説の各々の各単語には、入力音声における当該単語の持続時間を特定するための情報が付されており、スコア算出手段は、デコード手段が出力する複数の仮説の各々に対し、各単語の一般化単語事後確率を算出するための一般化単語事後確率算出手段と、一般化単語事後確率算出手段により算出された一般化単語事後確率と、単語ラティス中の各単語の持続時間を特定するための情報との関数としてスコアを各単語に対し算出するための関数計算手段とを含む。

【 0 0 3 2 】

関数計算手段は、各単語のスコアを算出するにあたり、当該単語の持続時間と、当該単語の一般化単語事後確率との関数としてスコアを定める。一般化単語事後確率だけでなく、単語の持続時間を考慮することにより、認識結果の信頼度を、その持続時間に対応する重み付けをして、経路全体のスコアに反映することができる。

【 0 0 3 3 】

最適経路探索手段は、単語ラティス内の始点から終点までの経路のうち、当該経路上の単語の各々に対し算出されたスコアの和が最大となる経路の上の単語列を音声認識結果として出力するための最大スコア経路探索手段を含んでもよい。

【 0 0 3 4 】

さらに好ましくは、関数計算手段は、以下の式によりスコアを各単語に対して算出するための手段を含む。

【 0 0 3 5 】

【数 1】

$$A_n = T_n \log C_n$$

ただし、 A_n はある仮説中の n 番目の単語のスコア、 T_n は入力音声の中の当該単語の持続時間、 C_n は当該単語に対して一般化単語事後確率算出手段により算出された一般化単語事後確率を、それぞれ示す。

【 0 0 3 6 】

本発明の第 2 の局面に係るコンピュータプログラムは、コンピュータにより実行されると、当該コンピュータを、上記したいずれかの音声認識システムとして動作させるものである。従って、コンピュータでこのコンピュータプログラムを実行させることにより、上記した音声認識システムと同様の効果を得ることができる。

【発明を実施するための最良の形態】

【 0 0 3 7 】

〔導入〕

本実施の形態では、上記した仮説の統合において、特許文献 1 で用いられた尤度に代えて、特許文献 2 で提案された一般化単語事後確率 (Generalized Word Posterior Probability: GWPP) を用いる。GWPP とは、音声認識の結果得られた単語ごとに、その音声認識結果の信頼度を示す尺度と考えられる。以下、GWPP について説明する。

【 0 0 3 8 】

仮説の統合という問題は、最終的には音声認識により認識された各単語を受け入れるか拒絶するかを判定することに帰着する。特許文献 2 は、この問題を、注目単語の位置の特定という考え方を導入することで解決している。この場合、注目単語以外の単語 (非注目単語) については、互いに区別せずいずれも単にそれぞれの場所を占めるだけのものとして取り扱って、注目単語の事後確率が算出される。

【 0 0 3 9 】

このように注目単語 / 非注目単語という二分法を採用することにより、動的計画法に基づく文字列のアライメント等の複雑な処理を行なう必要が回避できる。

【 0 0 4 0 】

以下、特許文献 2 の記載のうち、GWPP についてまとめる。まず、以下の概念を導入する。それらは、(1) 音声認識結果の単語ラティス又は N - ベストリスト中における、注目単語の位置決定を行なうための、仮説 (候補) となる文字列の探索空間の削減、(2) ある候補単語の複数の出現個所における事後確率をグループ化する際の時間的制約の緩和、及び (3) 音響モデル及び言語モデルによる寄与に対する適切なウェイト付け、である。

【 0 0 4 1 】

- 文字列と単語の事後確率 -

HMM を用いる音声認識装置では、所与の音響観測データ $x_1^T = x_1, \dots, x_T$ に対する、最適な単語シーケンス $w_1^{M*} = w_1^*, \dots, w_M^*$ を、以下に示すように、可能な全ての単語

10

20

30

40

50

語シーケンスからなる空間を探索して、最大事後確率 (MAP) を与えるものとして求める。

【 0 0 4 2 】

【 数 2 】

$$w_1^{M*} = \arg \max_{\{M, w_1^M\}} p(w_1^M | x_1^T) \quad (1)$$

$$= \arg \max_{\{M, w_1^M\}} \frac{p(x_1^T | w_1^M) p(w_1^M)}{p(x_1^T)} \quad (2)$$

$$= \arg \max_{\{M, w_1^M\}} p(x_1^T | w_1^M) p(w_1^M) \quad (3)$$

10

ただし、 $p(x_1^T | w_1^M)$ は音響モデルの確率、 $p(w_1^M)$ は言語モデルによる確率、 $p(x_1^T)$ は音響の観測確率である。

【 0 0 4 3 】

トレーニング環境とテスト環境、話者、雑音等の相違により「最適な」単語シーケンスであっても誤りを含むことがある。そこで、数学的に扱いやすく、かつ統計的に好ましい何らかの信頼度尺度を採用すべきである。

【 0 0 4 4 】

単語列の事後確率 $p(w_1^M | x_1^T)$ は、観測された音響 x_1^T に対し、認識された単語列 w_1^M の尤度を測るものであるが、これは対応する時間的セグメンテーション

20

【 0 0 4 5 】

【 数 3 】

$$[w; s, t]_1^M = [w_1; s_1, t_1] \cdots [w_M; s_M, t_M] \quad (4)$$

を仮定することで算出される。ただし、 s 及び t は単語 w の始点及び終点の時刻を示し、 $s_1 = 1$ 、 $t_M = T$ であり、 $1 \leq m \leq M-1$ の m に対し $t_m + 1 = s_{m+1}$ である。

【 0 0 4 6 】

これを用いて、式 (2) を次のように書き換えることができる。

30

【 0 0 4 7 】

【 数 4 】

$$p(w_1^M | x_1^T) = \frac{p(x_1^T | [w; s, t]_1^M) \cdot p(w_1^M)}{p(x_1^T)} \quad (5)$$

$$= \frac{\prod_{m=1}^M p(x_{s_m}^{t_m} | w_m) \cdot p(w_m | w_1^{m-1})}{p(x_1^T)} \quad (6)$$

認識された単語列の全体の信頼性を測るためには、この単語列事後確率 $p(w_1^M | x_1^T)$ を採用するのが自然である。

40

【 0 0 4 8 】

単語の信頼性を測るために適切な信頼度尺度は、単語事後確率 $p([w_m; s_m, t_m] | x_1^T)$ である。これは特定の単語を含む単語列の事後確率を全て合計することにより算出される。

【 0 0 4 9 】

【数 5】

$$\begin{aligned}
 & p([w; s, t] | x_1^T) \\
 &= \sum_{\substack{M, [w; s, t]_1^M \\ \exists n, 1 \leq n \leq M \\ [w_n; s_n, t_n] = [w; s, t]}} \frac{\prod_{m=1}^M p(x_{s_m}^{t_m} | w_m) p(w_m | w_1^{m-1})}{p(x_1^T)} \quad (7)
 \end{aligned}$$

この単語事後確率を実際に有効な信頼度尺度として用いるためには、さらにいくつかの問題を解決する必要がある。

10

【0050】

[単語事後確率の修正]

- 考慮すべき仮説数 -

大語彙の連続音声認識装置 (L V C S R) においては、可能な単語列の探索空間は膨大である。しかし、各単語列の事後確率の値には大きな相違があり、比較的低い尤度の単語列については刈込みしても差し支えない。このようにして得た、単語列の仮説の部分集合のみを用いて単語ラティス / グラフ又は N - ベスト単語列リストを得ることができる。以下の実施の形態では、そのように部分集合を用いて得た単語ラティス / グラフを使用するものとする。

【0051】

20

- 仮説内の単語の時間的なレジストレーション -

単語の時間的位置決め (レジストレーション) を $[w; s, t]$ で表わす。別々の仮説中にある同一の単語が出現する場合でも、その位置は仮説によって多少異なることがあり得る。自動音声認識 (A S R) の最終的目標は発話中の単語からなる内容を認識することであるから、厳密な時間的制約を多少緩和することにする。ここでは、ある単語がある単語列中において出現する期間が、基準となる単語の期間 $[s, t]$ と重なっており (オーバーラップしている)、かつその単語が基準となる単語と一致しているような単語を検索し、それら単語をその基準となる単語の事後確率の計算に含める。その結果式 (7) は以下のように書き換えられる。

【0052】

30

【数 6】

$$\begin{aligned}
 & p([w; s, t] | x_1^T) \\
 &= \sum_{\substack{M, [w; s, t]_1^M \\ \exists n, 1 \leq n \leq M \\ w = w_n \\ [s, t] \cap [s_n, t_n] \neq \emptyset}} \frac{\prod_{m=1}^M p(x_{s_m}^{t_m} | w_m) p(w_m | w_1^{m-1})}{p(x_1^T)} \quad (8)
 \end{aligned}$$

- 音響尤度と言語尤度との比重 -

40

本実施の形態では、音響尤度と言語尤度とには、それぞれ α 及び β で示されるウェイトによって指数的なウェイト付けがなされる。式 (8) にこれを適用すると次式となる。

【0053】

【数 7】

$$\begin{aligned}
 & p([w; s, t] | x_1^T) \\
 &= \sum_{\substack{M, [w; s, t]_1^M \\ \exists n, 1 \leq n \leq M \\ w = w_n \\ [s, t] \cap [s_n, t_n] \neq \emptyset}} \frac{\prod_{m=1}^M p^\alpha(x_{s_m}^{t_m} | w_m) p^\beta(w_m | w_1^M)}{p(x_1^T)} \quad (9)
 \end{aligned}$$

[注目単語の抽出]

10

ここで、本実施の形態に係る単語抽出方式により抽出された注目単語の受入 / 拒否について検討する。図 1 に本実施の形態で使用する、音声認識の結果得られる単語ラティス / グラフの例を示す。

【 0 0 5 4 】

図 1 を参照して、本実施例で使用する単語ラティス / グラフ 2 0 は、従来のものと異なり、注目単語（「 w 」で示す。）以外の単語については個々の単語ラベルを付さず、いずれも単に「 * 」というラベルを付してあるだけである。

【 0 0 5 5 】

この単語 w の出現個所の各々に対し、前方 - 後方アルゴリズムを用いて単語事後確率を効率的に計算できる。その後、この特定の単語 w（たとえば単語 3 0、3 2、3 4）を通るパスの全てについての尤度を合計し、その合計をこの単語ラティス / グラフ 2 0 内の全ての経路の尤度の合計で除算し正規化することによって、前述の G W P P が算出できる。さらにこの際、単語の時間的レジストレーション（単語開始及び終了時刻の一致）の条件を緩和する。すなわち、各経路の単語 w の期間が正確に一致する必要はなく、時間的にオーバーラップしているものの事後確率の合計を計算する。

20

【 0 0 5 6 】

[手法の概略]

雑音環境が頻繁に変動する状況では、音響モデルを高速に雑音環境に適応させることが可能でなければならない。以下に述べる本発明の一実施の形態では、高速な雑音環境適応として、非特許文献 1 4 において提案されている雑音 G M M の混合音適応化による H M M 合成法を用いる。

30

【 0 0 5 7 】

図 2 ~ 図 4 を参照して、この手法の概略について説明する。図 2 を参照して、あらかじめ準備した種々の雑音からなる雑音 D B 1 0 0 から、個々の雑音を混合成分とする雑音 G M M 1 0 2 と、個々の雑音に対して別々に適応化された雑音重畳音声用 H M M 1 0 4, 1 0 6, ... とを推定する。次に図 3 に示すように、短時間の未知雑音 1 1 0 を用いて雑音 G M M 1 0 2 の混合ウェイト $W_{N1}, W_{N2}, \dots$ のみを推定する。そして、図 4 に示すように、この混合ウェイト $W_{N1}, W_{N2}, \dots$ を用いて、雑音重畳音声用 H M M 1 0 4, 1 0 6, ... を状態レベルで複数混合化する。例えば H M M 1 0 4 の状態 1 2 0 と、H M M 1 0 6 の状態 1 2 2 とに対して、それぞれのガウス混合分布に対し図 3 に示すステップにより計算された混合ウェイトを乗算して足し合わせて状態出力確率分布 1 2 4 を算出し、雑音適応された H M M の状態 1 2 6 の状態出力確率分布とする。

40

【 0 0 5 8 】

図 2 ~ 図 4 において N_i は第 i 番目の雑音、 λ_i は第 i 番目の雑音に対する雑音重畳音声用 H M M を表す。 P_{Ni} と w_{Ni} は雑音 G M M における第 i 番目の雑音の分布とその分布に対する混合ウェイトとをそれぞれ示す。さらに $w_{\lambda_{ij}}$ と $p_{\lambda_{ij}}$ は第 i 番目の雑音用の雑音重畳音声用 H M M における第 j 番目の混合分布 N の分岐確率と混合成分とを表す。

【 0 0 5 9 】

この手法の利点として、適応の計算時間が G M M の混合ウェイトの推定時間のみであり大変高速である点と、雑音適応された H M M が複数の雑音環境の分布を含んでおり、単一

50

の雑音から推定された音響モデルよりも雑音の短時間の変動に対する頑健性が高い点とを挙げることができる。

【 0 0 6 0 】

上記した混合重み適応化によるHMM合成法を用いる場合、音響特徴量としてはMFCCを用いることが考えられる。しかし、MFCCのみでは認識精度を高めることが難しいことが実験的に判明した。そこで本実施の形態では、MFCCとは異なる音響特徴量を用いた音声認識を行ない、その結果とMFCCによる音声認識の結果とを統合することを考える。本実施の形態では、雑音の変動に対して頑健な特徴量として非特許文献4において提案されたDMFCC特徴量を用いることとする。以下、DMFCC特徴について述べる。なお、以下の処理では、音声データは所定サンプリング周波数及び所定窓長でサンプリングしたフレームとして準備されているものとする。

10

【 0 0 6 1 】

DMFCC特徴量は、式(10)に示すDPS(differential power spectrum)を基礎とする特徴量である。式(10)中の $Y(i, k)$ は、第 i 番目のフレームにおける第 k 番目のパワースペクトラム係数を表す。同様に $D(i, k)$ は第 i 番目のフレームにおける第 k 番目のDPS係数を表す。DMFCC特徴量は、このDPS係数に対してDCT(discrete cosine transform)を行なうことにより抽出される。

【 0 0 6 2 】

$$D(i, k) = |Y(i, k) - Y(i, k+1)| \quad (10)$$

20

有声母音等のピッチを含む音声から抽出されたパワースペクトラムは、基本周波数の高調波の影響によって櫛型の形状を持つ。このようなパワースペクトラムからDPS係数を計算した場合、隣り合うパワースペクトラム係数間の差が大きいため、DPS係数の値も同様に大きなパワーとして計算される。一方、雑音等の特徴を持たない波形のパワースペクトラムから計算されるDPS係数は、隣り合うパワースペクトラム係数間の差が小さいため、DPS係数の値も小さくなると考えられる。雑音重畳音声のパワースペクトラムを無雑音音声のパワーと雑音のパワーの和であると仮定した場合、DPS係数を計算することによって、音声と比較してなだらかに変化する雑音のパワー成分を減衰させることができると考えられる。

【 0 0 6 3 】

30

本実施の形態では、上述のようにMFCC特徴量とDMFCC特徴量とを用いて、パラレルにデコーディングを行ない、得られた仮説の統合による音声認識精度の改善を試みている。

【 0 0 6 4 】

[構成]

図5に、本実施の形態に係る音声認識システム130の概略ブロック図を示す。図5を参照して、このシステム130は、初期HMM150と、雑音データベース(DB)152と、雑音が重畳された学習データ153とから、パラレルに音声をデコードするためのMFCC・HMM群156及びDMFCC・HMM群158を作成するためのHMM作成部154と、HMM作成部154により作成されたMFCC・HMM群156及びDMFCC・HMM群158を用いて、入力音声144に対する音声認識を行ない、音声認識結果146を出力するための認識処理部142とを含む。

40

【 0 0 6 5 】

図6はHMM作成部154のブロック図である。図6を参照して、HMM作成部154は、初期HMM150と雑音DB152とから、前述したPMC法を用いて雑音重畳音声用MFCC・HMM群156を作成するためのMFCC雑音重畳音声用HMM推定部170と、雑音重畳済みの学習データ153を用いて初期HMM150に対する学習を行なうことにより、雑音重畳音声用DMFCC・HMM群158を作成するためのDMFCC雑音重畳音声用HMM推定部172とを含む。

【 0 0 6 6 】

50

本実施の形態では、雑音DB152としては12種類の異なる雑音を用いる。学習データ153についても、無雑音学習データに上記したものと同種の雑音を重畳したものを用いる。なお、雑音の重畳に際しては、10dB、20dB及び30dBの三種のSNRを用いている。初期HMM150としては、無雑音音響モデルとして学習済みのものを準備する。もちろん、後述する実験におけるように、重畳する雑音の種類及びそのSNRについてはこれに限定されるわけではない。

【0067】

MFC C雑音重畳音声用HMM推定部170は、従来技術の項で説明した通りのPMC法を用いて各雑音に対応する雑音重畳音声用HMMを推定する機能を持つ。同様にDMF C C雑音重畳音声用HMM推定部172は、学習データ153を用いて最尤推定を行なうことにより雑音重畳音声用DMF C C・HMM群158の学習を行なう。DMF C C特徴量に対しては、MFC C特徴量と異なりPMC法が適用できないためである。

10

【0068】

図7に、MFC C雑音重畳音声用HMM推定部170による雑音重畳音声用MFC C・HMM群156の概念について示す。図7を参照して、MFC C用の初期HMM180は、無雑音通常発声用MFC C・HMM190と、無雑音言直し発話用MFC C・HMM192とを含む。本実施の形態では、発話スタイルの変動への対応としてシステムへの言直し時に頻繁に観測される音節強調発話に対する頑健性の改善を試みている。言直し発話用のHMMはこのためのものである。

【0069】

20

音声認識ソフトウェアが認識誤りを起こした場合、そのソフトウェアの使用者はもう一度同じ発声を繰返さなければならない。このような言直し発話では、母音の後に短時間のポーズが挿入される等、通常発声とは異なる音響的特徴を持つことが報告されている。この言直し発話を頑健に認識するため、図18に示すような構造を持つ音響モデル440が提案されている。図18を参照して、この母音モデルは、母音の後に短時間ポーズを挿入するため、例えばt - a + s i lの状態パス(図18において、「t - a + k」等の表記は、先行音素が/t/、後続音素が/k/、当該音素が/a/の環境依存音素を表す。「s i l」は無音状態を表わす。)及び、その母音モデルの後にポーズ状態を追加した状態パスの合計三つの成分を有するマルチパス音響モデルの構造を持つ。さらに、このモデルでは、子音モデルの前に短時間ポーズの挿入を許すため、通常の子音モデルに加えてs i l - k + iの状態パスへの遷移が追加されている。このような音響モデルを用いることにより、通常発声の音声以外にも言直しや音節強調発声等の音声を頑健に認識することが可能となる。

30

【0070】

再び図7を参照して、雑音DB152は、本実施の形態では12種類の雑音データ200, 202, ..., 206を含む。MFC C雑音重畳音声用HMM推定部170はこれら12種類の雑音の各々について、3種類のSNR(10dB、20dB、及び30dB)ごとにPMCを用いて初期HMM180を適応化することにより、雑音重畳音声用MFC C・HMM群156を生成する。

【0071】

40

生成される雑音重畳音声用MFC C・HMM群156は、男声通常発声用MFC C・HMM群210と、男声言直し発話用MFC C・HMM群212と、女声通常発声用MFC C・HMM群214と、女声言直し発話用MFC C・HMM群216と、通常発声用無雑音MFC C・HMM215と、言直し発話用無雑音MFC C・HMM217とを含む。すなわち本実施の形態では、雑音重畳音声用MFC C・HMM群156は、男声女声、12種類の雑音、3種類のSNR、及び通常発声、言直し発話用の、 $2 \times 12 \times 3 \times 2 = 144$ 種類と通常発声用及び言直し発話用の無雑音音声用モデルの計146種類のHMMを含む。もちろん、条件によりこの個数が変化することは言うまでもない。

【0072】

図8に、MFC C雑音重畳音声用HMM推定部170により作成される音響モデルが、

50

音響空間 270 中に占める領域を模式的に示す。図 8 に示すのは、12 個の音響モデルに対応する領域 280 ~ 302 のみである。しかし、上述したように作成される音響モデルは 146 種類であるので、音響空間 270 にはこれら領域 280 ~ 302 と同様のものが合計で 146 個作成されることになる。

【0073】

図 9 に、DMFCC 雑音重畳音声用 HMM 推定部 172 による雑音重畳音声用 DMFCC・HMM 群 158 の作成を概念的に示す。図 9 を参照して、初期 DMFCC・HMM 182 は、無雑音通常発声用 DMFCC・HMM 230 及び無雑音言直し発話用 DMFCC・HMM 232 を含む。

【0074】

また雑音重畳学習データ 153 は、前述した 12 種類の雑音を、前述した 3 種類の SNR で学習データに重畳したものであり、 $3 \times 12 = 42$ 種類の雑音重畳学習データ 240 ~ 246 を含む。DMFCC 雑音重畳音声用 HMM 推定部 172 は、無雑音通常発声用 DMFCC・HMM 230 及び無雑音言直し発話用 DMFCC・HMM 232 に対し、上記した雑音重畳学習データ 153 を用いて学習を行なうことにより、男声通常発声用 DMFCC・HMM 群 250、男声言直し発話用 DMFCC・HMM 群 252、女声通常発声用 DMFCC・HMM 群 254、女声言直し発話用 DMFCC・HMM 群 256 と、通常発声用無雑音 DMFCC・HMM 255 と、言直し発話用無雑音 DMFCC・HMM 257 とを生成する。

【0075】

例えば男声通常発声用 DMFCC・HMM 群 250 は、各種類及び各 SNR の雑音重畳学習データに対して学習した結果得られた、複数個の男声雑音重畳通常発声用 DMFCC・HMM 260、262、...、266 を含む。他の DMFCC・HMM 群 252、254、256 も、男声か女声か、通常発声用モデルか言直し発話用モデルかを除き同様の構成である。

【0076】

本実施の形態では、雑音重畳音声用 DMFCC・HMM 群 158 は雑音重畳音声用 MFCC・HMM 群 156 と同様の構成となっている。しかし、当業者であれば容易に理解できるように、MFCC を用いる音声認識と、DMFCC を用いる音声認識とで同様の構成をとる必要は全くない。それぞれ別々のデータに基づき HMM を作成してもよい。最終的に作成される HMM の数が等しくなる必要もない。

【0077】

図 10 は、図 5 に示す認識処理部 142 の詳細な構造を示すブロック図である。図 10 を参照して、認識処理部 142 は、入力音声 144 に対し MFCC・HMM 群を用いて音声認識を行なう MFCC 処理部 310 と、入力音声 144 に対し DMFCC・HMM 群を用いた音声認識を行ない認識結果を出力するための DMFCC 処理部 312 と、MFCC 処理部 310 及び DMFCC 処理部 312 の出力する仮説を統合して出力するための仮説統合部 314 とを含む。

【0078】

図 11 は MFCC 処理部 310 のより詳細なブロック図である。図 11 を参照して MFCC 処理部 310 は、入力音声 144 から MFCC パラメータを音響特徴量として算出するための MFCC 算出部 320 と、MFCC 算出部 320 から出力される MFCC パラメータに対し、MFCC・HMM 群を用いて認識処理を行ない、HMM ごとに認識結果を出力するための MFCC 通常発声認識処理部 322 と、MFCC 算出部 320 から与えられる MFCC パラメータに対し、言直し発話用 HMM を用いて認識処理を行ない、HMM ごとに認識結果を出力するための MFCC 言直し発話認識処理部 324 と、MFCC 通常発声認識処理部 322 及び MFCC 言直し発話認識処理部 324 の出力のうち、尤度が最も高いものを選択して出力するための最尤選択部 326 とを含む。

【0079】

図 12 は、DMFCC 処理部 312 のより詳細なブロック図である。図 12 を参照して

10

20

30

40

50

DMFCC処理部312は、入力音声144から音響特徴量としてDMFCCパラメータを算出するためのDMFCC算出部330と、DMFCC算出部330から与えられるDMFCCパラメータに対しDMFCC通常発声用HMM群を用いて認識処理を行ない、認識結果をHMMごとに出力するためのDMFCC通常発声認識処理部332と、DMFCC算出部330からDMFCCパラメータを受取り、言直し発話用DMFCC・HMM群を用いて認識処理を行ない、HMMごとに認識結果を出力するためのDMFCC言直し発話認識処理部334と、DMFCC通常発声認識処理部332及びDMFCC言直し発話認識処理部334から出力される認識結果のうち、尤度が最も高いものを選択して出力するための最尤選択部336とを含む。

【0080】

10

図11及び図12を参照してわかるように、本実施の形態では、MFCC処理部310及びDMFCC処理部312の構造は互いに平行である。使用する音響特徴量がMFCCかDMFCCかによる差異があるにすぎない。従って以下では、MFCC処理部310の構造の詳細についてのみ説明する。

【0081】

図13は図11に示すMFCC通常発声認識処理部322のより詳細なブロック図である。図13を参照して、MFCC通常発声認識処理部322は、MFCC算出部320から与えられるMFCCパラメータに基づき、男声通常発声用MFCC・HMM群210及び女声通常発声用MFCC・HMM群214に対する雑音GMMの混合重み適応化によるHMM合成を重畳された雑音のSNRごとに行ない、男声通常発声用適応化MFCC・HMM群354及び女声通常発声用適応化MFCC・HMM群352を生成するための雑音適応化処理部350と、男声通常発声用適応化MFCC・HMM群354を用いて、入力されるMFCCパラメータに対するデコードを行なうことにより、適応化されたHMMごとにデコード結果を出力するためのMFCC男声通常発声デコーダ部358と、入力されるMFCCパラメータに対し女声通常発声用適応化MFCC・HMM群を用いてデコードし、HMMごとにデコード結果を出力するためのMFCC女声通常発声デコーダ部356とを含む。

20

【0082】

男声通常発声用適応化MFCC・HMM群354及び女声通常発声用適応化MFCC・HMM群352はそれぞれ、3種類のSNRごとに一つ、合計三個のHMMを含む。デコードには無雑音HMMも使用するので、デコーダ部356及び358はそれぞれデコード結果を4つずつ出力する。その結果、MFCC通常発声認識処理部322全体としては8つのデコード結果を出力する。

30

【0083】

ここで、図13に示す雑音適応化処理部350の処理について図16及び図17を参照して説明する。図16を参照して、雑音適応化処理部350は、入力されるMFCCパラメータに基づき、音響空間270中における入力音声に対応する領域420を推定する。そしてこの領域420と、予め求められている各種の雑音が占める領域280～302との距離を算出する。そして、距離が最も近いものを所定個数（本実施の形態では4つ）だけ選択する。図16の例で示せば領域290、292、296及び298により示される雑音が、入力される音声の音響空間中の領域420に最も近い。従って、この4つの雑音に対応する音響モデルが採用される。

40

【0084】

続いて図17を参照して、これら4つの領域290、292、296及び298に対応するHMMのガウス混合分布の重みを計算し、加算することにより、入力される音声の音響空間270中における領域420をカバーするような音響モデルをHMMの形で算出する。この音響モデルを用いて入力音声に対するデコードを行なう。このように各雑音に対する音響モデル自体は変化させず今後のための重みのみを計算して音声認識用のHMMの適応化を行なえばよい。そのため適用の計算時間が短く、大変高速に適応化を行なうことができる。さらに、適応化されたHMMが複数の雑音環境の分布を含んでいる。従って単

50

数の雑音から推定された音響モデルを用いた場合よりも、雑音の短時間の変動に対する頑健性がより高くなるという利点がある。

【 0 0 8 5 】

図 1 4 は M F C C 言直し発話認識処理部 3 2 4 の構成を示す。M F C C 言直し発話認識処理部 3 2 4 は、入力される M F C C パラメータを用いて、男声言直し発話用 M F C C ・ H M M 群 2 1 2 及び女声言直し発話用 M F C C ・ H M M 群 2 1 6 に対し雑音 G M M の混合重み適応化による H M M 合成法を重畳された雑音の S N R ごとに行ない、男声、女声及び S N R ごとに適応化された H M M を出力することにより、男声言直し発話用適応化 M F C C ・ H M M 群 3 7 4 及び女声言直し発話用適応化 M F C C ・ H M M 群 3 7 2 を出力するための雑音適応化処理部 3 7 0 と、与えられる M F C C パラメータを、女声言直し発話用適応化 M F C C ・ H M M 群 3 7 2 を用いてデコードし、H M M ごとに出力するための M F C C 女声言直し発話デコーダ部 3 7 6 と、入力される M F C C パラメータを男声言直し発話用適応化 M F C C ・ H M M 群 3 7 4 を用いてデコードし、H M M ごとにデコード結果を出力するための M F C C 男声通常発声デコーダ部 3 7 8 とを含む。

10

【 0 0 8 6 】

女声言直し発話用適応化 M F C C ・ H M M 群 3 7 2 は、S N R ごとに合成される三つの H M M を含む。男声言直し発話用適応化 M F C C ・ H M M 群 3 7 4 も同様に、S N R ごとの三つの H M M を含む。また、デコードには無雑音 H M M も使用される。従って、デコーダ部 3 7 6 及び 3 7 8 はそれぞれ 4 つずつのデコード結果を出力する。その結果 M F C C 言直し発話認識処理部 3 2 4 の出力は 8 つとなる。

20

【 0 0 8 7 】

図 1 3 及び図 1 4 を参照して明らかなように、本実施の形態においては、M F C C 通常発声認識処理部 3 2 2 と M F C C 言直し発話認識処理部 3 2 4 との構成はパラレルである。従って以下では M F C C 通常発声認識処理部 3 2 2 の詳細な構造のみを説明する。また図 1 3 及びこれ以前の説明から明らかなように、M F C C 女声通常発声デコーダ部 3 5 6 及び M F C C 男声通常発声デコーダ部 3 5 8 の構成も互いにパラレルである。従って以下では女声についてのみ M F C C 通常発声認識処理部 3 2 2 の詳細な構成を説明する。

【 0 0 8 8 】

図 1 5 は、M F C C 女声通常発声デコーダ部 3 5 6 及び女声通常発声用適応化 M F C C ・ H M M 群 3 5 2 の詳細な構成を示す。図 1 5 を参照して、女声通常発声用適応化 M F C C ・ H M M 群 3 5 2 は、無雑音 H M M 4 0 2 と、それぞれ 1 0 d B、2 0 d B、及び 3 0 d B の S N R で雑音が重畳された雑音重畳 H M M から合成された 1 0 d B 雑音 H M M 4 0 4、2 0 d B 雑音 H M M 4 0 6、及び 3 0 d B 雑音 H M M 4 0 8 とを含む。

30

【 0 0 8 9 】

M F C C 女声通常発声デコーダ部 3 5 6 は、入力される M F C C パラメータを、無雑音 H M M 4 0 2、1 0 d B 雑音 H M M 4 0 4、2 0 d B 雑音 H M M 4 0 6、及び 3 0 d B 雑音 H M M 4 0 8 をそれぞれ用いてデコードし、デコード結果を出力するためのデコーダ 3 9 0、3 9 2、3 9 4、及び 3 9 6 を含む。

【 0 0 9 0 】

図 1 9 に、図 1 0 に示す仮説統合部 3 1 4 のより詳細な構成を示す。図 1 0 に示す M F C C 処理部 3 1 0 及び D M F C C 処理部 3 1 2 からは複数の仮説が仮説統合部 3 1 4 に与えられる。仮説統合部 3 1 4 は、これら複数の仮説を単語単位で統合する。その原理について図 2 0 ~ 図 2 2 を参照して説明する。なお、以下の説明に用いられる G W P P については図 1 を用いて冒頭に説明したとおりである。

40

【 0 0 9 1 】

複数の音声認識デコーダから得られた仮説が互いに相補的である場合、それぞれの仮説の正しい部分を抽出して組合わせることにより、より正しい単語列が得られる可能性がある。ここで「相補的」とは、あるデコーダの認識結果の前半は正しいが後半は間違いであったとしても、別のデコーダの認識結果の後半部分が正しいならば、それぞれの正しい部分をつなぎあわせることによりその認識誤りを補償することができるという意味である。

50

【 0 0 9 2 】

図 2 0 を参照して、二つの仮説 4 7 0 及び 4 7 2 が得られたものとする。仮説 4 7 0 の前半部分は誤っているが後半部分は正しい認識結果である。一方、仮説 4 7 2 については、前半の認識結果は正しいが後半は誤りである。従って仮説 4 7 2 の前半部分と仮説 4 7 0 の後半部分とをつなぎ合わせることで、正しい結果が得られるはずである。

【 0 0 9 3 】

図 2 1 を参照して、上記した結果を得るために、まず図 2 1 に示されるような単語ラティス 4 8 0 を、与えられた二つの仮説から再構成する。この再構成では、個々の単語の開始及び終了時間情報を用いる。

【 0 0 9 4 】

続いて図 2 2 に示されるように、この単語ラティス 4 8 0 において音響尤度と言語尤度、並びに言語モデルから各仮説の単語ごとに G W P P を算出する。この G W P P の関数として各単語に対してスコアを付与する。そして、単語ラティス 4 8 0 内の、始点と終点とを結ぶ単語列経路のうち、単語列全体としてのスコアが最も大きくなるような単語列 4 8 2 を再探索する。特許文献 2 にも開示されているとおり、仮説のうちでも音声認識の信頼性の高い部分の G W P P は高く、信頼性が低い部分の G W P P は低くなっている。従って、このような再探索を行なうことにより二つの仮説を統合して正しい結果を得ることができる可能性が高くなる。

【 0 0 9 5 】

ただし、単語ラティスの再構成と、G W P P の算出との順序はこれと逆でもよい。G W P P の算出は、仮説ごとに行なうためである。

【 0 0 9 6 】

なお本実施の形態では、M F C C と D M F C C 特徴量から得られた仮説に対する仮説統合を仮説統合部 3 1 4 で行なっている。この仮説統合の際には、言語モデルを用いた尤度計算も行なう。この場合、音響モデルの尤度計算と言語モデルによる尤度計算との間でのウェイト付けを考慮しなければならない。

【 0 0 9 7 】

図 1 9 を参照して、仮説統合部 3 1 4 は、上記したような機能を実現するために以下の各処理部を含む。すなわち仮説統合部 3 1 4 は、M F C C 及び D M F C C の仮説を保持し、後述する様に各単語に対し算出される G W P P を付与するように各仮説を更新するための仮説更新部 4 5 2 と、仮説更新部 4 5 2 が保持する仮説の各々の各単語について、G W P P を算出するために、単語ラティス中で期間がオーバーラップする、同じ単語を検索するための対象単語検索部 4 5 4 と、対象単語検索部 4 5 4 により検索された単語群に対し、前述した算出方法によりその G W P P を算出するための G W P P 算出部 4 5 6 と、G W P P 算出部 4 5 6 による G W P P 算出の際に参照される、統計的言語モデルを記憶するための言語モデル記憶部 4 6 0 と、G W P P 算出の際の言語モデルの尤度と、音響モデルの尤度とのウェイトを記憶するためのウェイト記憶部 4 5 8 とを含む。仮説更新部 4 5 2 は、G W P P 算出部 4 5 6 により算出された G W P P を用いて以下の式により算出されるスコア A_n を各仮説の単語に付して出力する機能を有する。

【 0 0 9 8 】

【数 8】

$$A_n = T_n \log C_n$$

ただし n は各仮説中における単語の順番を示し、 T_n はその単語の持続時間を示し、 C_n はその単語に付与された G W P P を示す。本実施の形態では、単語列全体のスコアは、こうして各単語に対して算出されたスコアの和として算出する。このようにスコアを G W P P だけでなく単語の持続時間とも関連付けることにより、G W P P にその単語の持続時間に相当する重みを付けることになり、仮説全体の信頼度に各単語の G W P P をよりよく反映できる。

【 0 0 9 9 】

なお、ウェイト記憶部 4 5 8 に記憶されるウェイトについては、特許文献 2 にウェイト

10

20

30

40

50

α 及び β として記載がある。特許文献 2 では、単独の音声認識部から出力される単語ラティスについてこの G W P P を適用した最適経路探索を行なっているが、本実施の形態でも特許文献 2 と同様のウェイトを用いる。

【 0 1 0 0 】

仮説統合部 3 1 4 はさらに、仮説更新部 4 5 2 が出力する二つの仮説から、個々の単語の開始及び終了時間情報、並びにスコアを用いて単語ラティス 4 8 0 (図 2 1 参照) を作成するための単語ラティス作成部 4 6 2 と、単語ラティス作成部 4 6 2 により作成された単語ラティスを記憶するための単語ラティス記憶部 4 6 4 と、単語ごとにスコアが付与された単語ラティスの中で、経路上の単語列の G W P P の値の和が最も高い経路 (最高スコア経路) を探索して、その経路に含まれる単語列を、G W P P とともに音声認識結果 1 4 6 として出力するための最高スコア経路探索部 4 6 6 とを含む。最高スコア経路探索部 4 6 6 の出力する単語列は、仮説統合部 3 1 4 に入力される二つの仮説を統合して得た、最も信頼性の高い仮説となる。

10

【 0 1 0 1 】

< コンピュータによる実現 >

上記した実施の形態に係る音声認識システム 1 3 0 は、コンピュータハードウェアと、当該コンピュータハードウェアの上で C P U により実行されるコンピュータプログラムとにより実現可能である。

【 0 1 0 2 】

図 2 3 に、特に仮説統合部 3 1 4 をコンピュータで実現するためのコンピュータプログラムのフローチャートを示す。図 2 3 を参照して、仮説統合部 3 1 4 に相当する処理は、二つの仮説を受取ると次の処理を行なう。これら二つの仮説は、いずれもノードとアークとからなるグラフ形式で与えられる。アークは仮説を構成する単語に相当し、ノードは単語と単語との接続部に相当する。各アークには、入力された音声信号に基づき、音声認識の際にその単語の開始事項と終了時刻とが付与されている。

20

【 0 1 0 3 】

最初に、ステップ 4 9 0 において、各仮説の各単語に対して G W P P を計算する。この際には、実施の形態の冒頭に説明したように、各単語ごとに、順に G W P P を計算し付与する。

【 0 1 0 4 】

ステップ 4 9 2 では、このようにして各単語に G W P P が付与された二つの仮説のうち、同一時間に生起している同一単語があれば、それらに対応するアークを一つのアークにまとめる。

30

【 0 1 0 5 】

このようにしてアークの統合について可能なものを全て行なった後、ステップ 4 9 4 では、結果として得られた単語ラティスの中で隣接するアーク間にノードを挿入する。より具体的には、一つのアークに付与されている単語の終了時刻と、次のアークに付与されている単語開始時刻とが 3 0 ミリ秒以内であるようなアーク対があれば、その間にノードを挿入する。このようにして、可能な限りノードを挿入する。この結果、各アークに G W P P が付与された単語ラティスが作成される。

40

【 0 1 0 6 】

ステップ 4 9 6 では、単語ラティスの各アークについて、付与されている G W P P と、開始及び終了時刻とに基づき、そのアークに対応する単語のスコア A_n が算出される。スコアを全てのアークに対して算出し付与することにより、スコア付きの、統合された単語ラティスが完成する。

【 0 1 0 7 】

最後に、ステップ 4 9 8 で、この単語ラティスの中で最高のスコアが得られる経路を、D P サーチにより探索し、得られた経路上の単語列を出力して処理を終了する。

【 0 1 0 8 】

図 2 4 は、この音声認識システム 1 3 0 を実現するコンピュータシステム 5 3 0 の外観

50

を示し、図 25 はコンピュータシステム 530 の内部構成を示す。

【0109】

図 24 を参照して、このコンピュータシステム 530 は、メモリドライブ 552 及び DVD ドライブ 550 を有するコンピュータ 540 と、キーボード 546 と、マウス 548 と、モニタ 542 と、音声入力に用いられるマイクロフォン 570 と、一対のスピーカ 572 とを含む。

【0110】

図 25 を参照して、コンピュータ 540 は、メモリドライブ 552、DVD ドライブ 550、マイクロフォン 570 及びスピーカ 572 に加えて、CPU (中央処理装置) 556 と、CPU 556、メモリドライブ 552 及び DVD ドライブ 550 に接続されたバス 566 と、ブートアッププログラム等を記憶する読出専用メモリ (ROM) 558 と、バス 566 に接続され、プログラム命令、システムプログラム、及び作業データ等を記憶するランダムアクセスメモリ (RAM) 560 と、バス 566 に接続され、マイクロフォン 570 及びスピーカ 572 を用いた音声処理を行なうサウンドボード 568 とを含む。

【0111】

ここでは示さないが、コンピュータ 540 はさらにローカルエリアネットワーク (LAN) への接続を提供するネットワークアダプタボードを含んでもよい。

【0112】

コンピュータシステム 530 に音声認識装置 130 としての動作を行なわせるためのコンピュータプログラムは、DVD ドライブ 550 又はメモリドライブ 552 に挿入される DVD 562 又は不揮発性メモリ 564 に記憶され、さらにハードディスク 554 に転送される。又は、プログラムは図示しないネットワークを通じてコンピュータ 540 に送信されハードディスク 554 に記憶されてもよい。プログラムは実行の際に RAM 560 にロードされる。DVD 562 から、不揮発性メモリ 564 から、又はネットワークを介して、直接に RAM 560 にプログラムをロードしてもよい。

【0113】

このプログラムは、コンピュータ 540 にこの実施の形態の音声認識装置 130 としての動作を行なわせる複数の命令を含む。この動作を行なわせるのに必要な基本的機能のいくつかはコンピュータ 540 上で動作するオペレーティングシステム (OS) もしくはサードパーティのプログラム、又はコンピュータ 540 にインストールされる各種ツールキットのモジュールにより提供される。従って、このプログラムはこの実施の形態のシステム及び方法を実現するのに必要な機能全てを必ずしも含まなくてよい。このプログラムは、命令のうち、所望の結果が得られるように制御されたやり方で適切な機能又は「ツール」を呼出すことにより、上記した音声認識装置 130 としての動作を実行する命令のみを含んでいればよい。コンピュータシステム 530 の動作は周知であるので、ここでは繰返さない。

【0114】

[動作]

上記した音声認識システム 130 は以下のように動作する。図 26 に、このシステムの動作の概略の流れについて示す。大きく分けて、このシステムは二つの動作局面を持つ。第一の局面は、雑音重畳音声用の HMM を準備するステップ 500 である。第二の局面は、このようにして準備された雑音重畳音声用の HMM と無雑音用の HMM とを用いて、入力される音声の認識を行なうステップ (502 ~ 508) である。

【0115】

ステップ 500 では、図 5 に示すような初期 HMM 150 と、雑音 DB 152 とを用いて、M F C C ・ HMM 群 156 が作成され、また雑音重畳学習データ 153 を用いて D M F C C ・ HMM 群 158 が作成される。

【0116】

このようにして、雑音重畳音声用の HMM 群が作成された後は、いつでもこの M F C C ・ HMM 群 156 及び D M F C C ・ HMM 群 158 を用いた音声認識を行なうことができ

10

20

30

40

50

る。図5に示す入力音声144が与えられると、その入力音声からMFCCパラメータ及びDMFCCパラメータが算出される(ステップ502)。それらを用いて、予め準備されたMFCC・HMM群156及びDMFCC・HMM群158のうち入力音声144の発話環境に最も類似した発話環境に対応する所定個数(本実施の形態では4個)のHMMがMFCC及びDMFCCのそれぞれについて選択される。これらHMMからMFCC及びDMFCCの各々について、雑音GMMの混合重み適応化によるHMMが合成される。合成されるHMMは、男声・女声、通常発声・言直し発話、及び4種類のSNR(10dB、20dB、30dB、無雑音)の組合わせの各々に対してであるから、全部で $2 \times 2 \times 4 = 16$ 通りである。

【0117】

10

続いてステップ504で発話入力があったか否かが判定される。発話入力があればステップ506に進むが、発話入力がなければ、再び重み推定502を行なう。本実施の形態では、発話入力があった場合には、その直前の1秒間の期間における雑音を用いて重み推定を行なっている。

【0118】

ステップ506では、合成されたHMMを用いた音声認識と、それら音声認識により得られた仮説を統合して最終的な単語ラティスを作成する処理とが行なわれる。この単語ラティスの各単語には、各単語の持続時間とGWPPとの関数として前述した式により算出されるスコアが付与される。単語ラティスの中で最も高いスコアを与える経路を探索し、その結果得られた経路上の単語列がステップ508で出力される。この後、重み推定502からの処理が繰返される。

20

【0119】

ステップ506での仮説統合部314の動作の詳細を説明する。図19を参照して、仮説更新部452は、MFCC処理部310及びDMFCC処理部312から与えられた二つの仮説を保持する。

【0120】

対象単語検索部454は、仮説更新部452に保持された二つの仮説の各々について、GWPPを算出する対象の単語を順番に取り出し、GWPP算出部456に与える。GWPP算出部456は、この単語について、仮説更新部452に保持された仮説の構造と、言語モデル記憶部460に記憶された統計的言語モデルと、ウェイト記憶部458に記憶された、GWPP算出時に音響モデルによる尤度と言語モデルによる尤度とにそれぞれ割当てられるウェイトとを使用して、GWPPを算出し仮説更新部452に与える。仮説更新部452は、与えられたGWPPと、仮説中の当該単語の持続時間に関する情報とに基づいて、当該単語のスコアを算出し付与する。

30

【0121】

二つの仮説の各単語に対し、仮説更新部452、対象単語検索部454、及びGWPP算出部456によってGWPPが算出されると、その二つの仮説は単語ラティス作成部462に与えられる。単語ラティス作成部462は、図23にフローチャートを示した処理によって単語ラティスを作成し、単語ラティス記憶部464に記憶させる。最高スコア経路探索部466は、単語ラティス記憶部464に記憶された単語ラティスの始点と終点とを結ぶ経路上で、経路上の単語のGWPPの値の和が最高となるものを探索し、該当する経路上の単語列を統合後の仮説である音声認識結果146として出力する。

40

【0122】

図27を参照して、上記実施の形態では、ある発話522に対しては、発話522の直前の雑音524を用いて合成されたHMMによる音声認識が行なわれる。同様に次の発話526に対しては、発話526の直前の雑音528により推定されたHMMを用いて音声認識が行なわれる。

【0123】

なお、上記した男声女声、MFCC及びDMFCC、通常発声及び言直し発話等の組合せは任意に選ぶことができる。MFCCサブシステム又はDMFCCサブシステムのいず

50

れか一方のみを用いるシステムも可能である。

【 0 1 2 4 】

また、上記した実施の形態では、最終的に二つの仮説を得て、それらを統合した。しかし本発明はそのような実施の形態には限定されない。例えば三つ以上の仮説を最終段階で統合するようにしてもよい。

【 0 1 2 5 】

[実験]

上記した実施の形態に係るシステムの有効性を検証するために、以下の実験を行なった。実験は、A U R O R A - 2 J タスクで行なった。このデータベースは学習及び試験のための、日本語の数字の連続発話コーパスを含んでいる。本件出願人により作成された A T R A S R のバージョン 3 . 3 をデコーダとして用いた。音響モデルを推定するためには、A U R O R A - 2 J 中のクリーン学習セットを使用した。このセットには、1 1 0 名の発話者（男性 5 5 名、女性 5 5 名）の 8 , 4 4 0 発話が含まれている。このトレーニングセットに、レストラン、街頭、空港、駅という 4 種類の雑音を 4 種類の S N R (2 0 , 1 5 , 1 0 及び 5 d B) で重畳した。数字発話と無音状態とを表 1 に示すような種々の H M M でモデリングした。その結果得られた音響モデルは、二つの特徴量 (M F C C , D M F C C) × 4 種類の S N R × 4 種類の雑音 = 3 2 通りである。各雑音種類ごとに、8 ガウス分布の雑音 G M M の学習を行なった。

【 0 1 2 6 】

【表 1】

	ステート数	混合分布数
数字発話	16	20
無音状態	3	36
短いポーズ	1	36

テストは A U R O R A - 2 J のテストセット B を用いて行なった。このセットでは、発話データに、学習データに重畳したものと異なる 4 種類の雑音（地下鉄、バブル、車内、展示会）を、前述した 5 d B から 2 0 d B の 4 種類の S N R と、無雑音との、合計 5 種類の S N R で重畳した。学習データとテストデータとのいずれに対しても、G . 7 1 2 フィルタを適用した。

【 0 1 2 7 】

M F C C の特徴量ベクトルには、1 0 ミリ秒フレーム間隔で 2 5 ミリ秒フレーム長のフレームにより算出した 1 2 個の M F C C 、 $\Delta p o w$ 、1 2 個の $\Delta M F C C$ 、 $\Delta \Delta p o w$ 、及び 1 2 個の $\Delta \Delta M F C C$ を用いた。D M F C C 特徴量も M F C C 特徴量と同じ構成のものを用いた。ケプストラム平均除去をいずれの特徴量にも適用した。これら特徴量をそれぞれ M F C C _ C M S 及び D M F C C _ C M S と呼ぶことにする。さらに、特徴量の抽出に先立ち、2 段ウィーナーフィルタリング (E T S I (欧州通信規格協会)) により配布されている A F E (E S 2 0 2 0 5 0 アナログフロントエンド)) を適用した。A F E により雑音抑制した M F C C 及び D M F C C の特徴量を、それぞれ M F C C _ A F E 及び D M F C C _ A F E と呼ぶことにする。

【 0 1 2 8 】

上記した条件で、本発明に係るシステムで使した M F C C サブシステム及び D M F C C サブシステム、本発明に係るシステム、並びに特許文献 1 に記載の尤度による仮説統合を採用したシステムにおいて、それぞれ C M S 及び A F E を適用した場合の単語認識精度を測定した。結果を表 2 に示す。

【 0 1 2 9 】

【表 2】

	+CMS	+AFE
MFCCサブシステム	89.21	91.24
DMFCCサブシステム	89.38	91.17
GWPPによる仮説統合	89.80	91.90
尤度による仮説統合	89.38	91.81

表 2 から明らかなように、M F C C サブシステム及び D M F C C サブシステム単独の場合よりも、仮説統合を行なった場合の方が高い精度を示す。しかも、本願実施の形態に係る G W P P による仮説統合を採用したシステムの精度は、特許文献 1 に開示された、尤度

10

【0130】

次に、雑音及び発話スタイルについて、本願実施の形態に係るシステムの頑健さがどの程度かを調べるため、通常発話と言直し発話とについての次のような実験を行なった。主な条件は以下の通りである。

【0131】

雑音に対する高速適応化のための雑音ごとの音響モデルを、A T R 旅行発話データベース（5 時間）の対話発話、音素バランス文の読上げ発話（2.5 時間）及び空港、地下街等 12 種類の雑音を使用してトレーニングした。M D L - S S S 手法を用いて、2,089 状態の状態共有構造が生成された。各状態は 5 個のガウス分布要素を有していた。全ての音響モデルに、種々の雑音を 10、20 及び 30 d B の S N R でそれぞれ重畳した。言直し発話用の音響モデルは通常発話の音響モデルから生成した。各分布のパラメータは同一としたが、H M M のトポロジーは色々であった。各音響モデルは発話者の性に依存したものである。従って、得られた音響モデルは、S N R の 3 レベル、雑音の 12 種類、M F C C 及び D M F C C という特徴量、発話者の性、及び発話スタイルに依存しており、合計 $3 \times 12 \times 2 \times 2 \times 2 = 288$ 通りである。

20

【0132】

各雑音タイプに対し、8 ガウス分布を持つ雑音 G M M の学習を行なった。テスト発話の各々の先頭の 500 ミリ秒の部分を使用して雑音適応により適応化音響モデルを作成した。M F C C 特徴量は、10 ミリ秒のフレーム間隔かつ 20 ミリ秒のフレーム長のフレームから得た 12 個の M F C C , 12 個の Δ M F C C 及び Δ p o w からなる。D M F C C 特徴量は、同様に 12 個の D M F C C , 12 個の Δ D M F C C 及び Δ p o w からなる。

30

【0133】

言語モデルとしては、マルチクラスの複合単語バイグラム及び単語トライグラムのものを用いた。言語モデルの単語数は合計 6.1 M 単語であり、レキシコンサイズは 34 k 単語であった。

【0134】

通常発話には B T E C コーパスのテストセット 1 を用いた（510 文、男性 4 人及び女性 6 人）。言直し発話には男性 2 名、女性 2 名の、文節ごとに強調した 40 個の文を収集した。テスト用の通常発話には 3 種類の雑音を 4 通りの S N R レベルで重畳した。言直し発話には 3 種類の雑音を 3 通りの S N R レベルで重畳した。

40

【0135】

通常発話用の音響モデルを有するシステムと、通常発話及び言直し発話用の音響モデルを持つシステムとの双方について、M F C C サブシステムのみ、D M F C C サブシステムのみ、G W P P を用いた仮説統合を行なうシステム（本実施の形態）、及び尤度による仮説統合を採用したシステム（非特許文献 1）によって音声認識性能を測定した。いずれのシステムにおいても、入力発話の最初の 500 ミリ秒を用いて雑音環境に対する音響モデルの適応化を行なった。表 3 に、通常発話に対する音声認識精度（%）を示す。なお表 3

50

において、「通常発話のみ」は通常発話用の音響モデルのみを用いたシステムであることを示し、「通常発話＋言直し発話」は、通常発話用の音響モデルと言直し用音響モデルとの双方を用いたシステムであることを示す。

【 0 1 3 6 】

【表 3】

音響モデル	通常発話のみ	通常発話＋言直し発話
MFCCサブシステム	87.30	86.40
DMFCCサブシステム	86.78	85.55
GWPPによる仮説統合	87.69	86.74
尤度による仮説統合	86.84	85.89

10

表 3 によれば、非特許文献 1 に記載の尤度による仮説統合を採用したシステムの場合、MFCCサブシステムのみの場合の精度より悪化しているという結果が得られた。それに対し、本実施の形態に係る、GWPPによる仮説統合を採用するシステムでは、いずれの場合も、どちらのサブシステムよりも高い認識精度が得られた。

【 0 1 3 7 】

言直し発話に対する実験では、通常発話用の音響モデルのみのシステムでは 10 % 程度の精度しか得られなかったが、本実施の形態のように言直し発話用の音響モデルをさらに追加して使用すると、単語精度として MFCC サブシステムのみ、DMFCC サブシステムのみの場合、種々の SNR の場合を平均してそれぞれ約 37 % 及び約 36 % という結果

20

【 0 1 3 8 】

以上のように本実施の形態の音声認識システム 130 では、雑音と発話スタイルの変動とに頑健で、特許文献 1 に開示のシステムよりもさらに精度の高い音声認識を実現することを目指した。本システムでは、雑音の変動に頑健な音響特徴量としての DMFCC、予め種々の雑音環境に適応化した HMM を用いて雑音 GMM の混合ウェイトから雑音適応 HMM を高速に生成する雑音適応手法、言直し発話に頑健な音響モデル、及び複数の仮説を GWPP を用いて統合する手法を用いた。その結果、種々の SNR で雑音を重畳した通常発声の評価データに対して、特許文献 1 に開示の尤度による仮説統合を採用したシステム

30

よりも高い単語認識精度が得られた。また、言直し発話等の発話スタイルの変動に対して、通常発声用音響モデルのみを用いた場合よりも高く、かつ特許文献 1 に開示のシステムよりも高い単語認識精度が得られた。

【 0 1 3 9 】

今回開示された実施の形態は単に例示であって、本発明が上記した実施の形態のみに制限されるわけではない。本発明の範囲は、発明の詳細な説明の記載を参酌した上で、特許請求の範囲の各請求項によって示され、そこに記載された文言と均等の意味及び範囲内のすべての変更を含む。

【図面の簡単な説明】

【 0 1 4 0 】

40

【図 1】本発明の第 1 の実施の形態に係る音声認識システムで採用した GWPP の算出原理を説明するための、単語ラティスの模式図である。

【図 2】雑音 GMM 及び雑音重畳音声 HMM の作成を説明するための図である。

【図 3】混合重みの推定を説明するための図である。

【図 4】適応化 HMM の生成を説明するための図である。

【図 5】本発明の一実施の形態に係る音声認識システムのブロック図である。

【図 6】HMM 作成部のより詳細なブロック図である。

【図 7】本発明の一実施の形態における雑音重畳音声用 MFCC・HMM 群の作成を説明するための図である。

【図 8】雑音 GMM の混合重み適応化において、PMC 法により準備される雑音 HMM を

50

説明するための図である。

【図 9】本発明の一実施の形態において雑音重畳音声用 D M F C C ・ H M M を作成する方法を説明するための図である。

【図 10】認識処理部のより詳細な構成を示すブロック図である。

【図 11】M F C C 処理部 3 1 0 の詳細な構成を示すブロック図である。

【図 12】D M F C C 処理部 3 1 2 の詳細な構成を示すブロック図である。

【図 13】M F C C 通常発声認識処理部の詳細な構成を示すブロック図である。

【図 14】M F C C 言直し発話認識処理部の詳細な構成を示すブロック図である。

【図 15】M F C C 女声通常発声デコーダ部 3 5 6 及び女声通常発声用適応化 M F C C ・ H M M 群 3 5 2 の詳細な構成を示すブロック図である。

10

【図 16】本実施の形態における入力音声の発話環境から、予め準備された雑音 H M M の発話環境までの距離を概念的に説明するための図である。

【図 17】入力音声の発話環境に類似した雑音を含む雑音 H M M から適応化 H M M を合成する概念を示す図である。

【図 18】言直し発話に頑健な音響モデルの構成を示す図である。

【図 19】仮説統合部の詳細な構成を示すブロック図である。

【図 20】仮説統合の経過を説明するための、二つの仮説を示す図である。

【図 21】仮説統合の過程で生成される単語ラティスを示す図である。

【図 22】G W P P による仮説統合の際に行なわれる、最も高い G W P P を示す単語列の探索を説明するための図である。

20

【図 23】図 1 9 に示す仮説統合部を実現するためのコンピュータプログラムのフローチャートである。

【図 24】本発明の一実施の形態に係る音声認識装置 1 3 0 を実現するコンピュータシステムの外観図である。

【図 25】図 2 4 に示すコンピュータのブロック図である。

【図 26】本発明の一実施の形態に係る音声認識システムの動作を説明するための図である。

【図 27】発話ごとの音声認識に用いられる雑音の位置を説明するための図である。

【図 28】P M C 法の概念を説明するための図である。

【符号の説明】

30

【 0 1 4 1 】

1 3 0 音声認識システム

1 4 2 認識処理部

1 4 6 音声認識結果

1 5 4 H M M 作成部

1 5 6 雑音重畳音声用 M F C C ・ H M M 群

1 5 8 雑音重畳音声用 D M F C C ・ H M M 群

1 9 0 無雑音通常発声用 M F C C ・ H M M

1 9 2 無雑音言直し発話用 M F C C ・ H M M

2 1 0 男声通常発声用 M F C C ・ H M M 群

40

2 1 2 男声言直し発話用 M F C C ・ H M M 群

2 1 4 女声通常発声用 M F C C ・ H M M 群

2 1 6 女声言直し発話用 M F C C ・ H M M 群

2 3 0 無雑音通常発声用 D M F C C ・ H M M

2 3 2 無雑音言直し発話用 D M F C C ・ H M M

2 5 0 男声通常発声用 D M F C C ・ H M M 群

2 5 2 男声言直し発話用 D M F C C ・ H M M 群

2 5 4 女声通常発声用 D M F C C ・ H M M 群

2 5 6 女声言直し発話用 D M F C C ・ H M M 群

3 1 0 M F C C 処理部

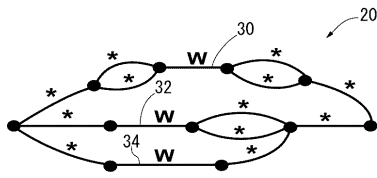
50

3 1 2 D M F C C 処理部
 3 1 4 仮説統合部
 3 2 0 M F C C 算出部
 3 2 2 M F C C 通常発声認識処理部
 3 2 4 M F C C 言直し発話認識処理部
 3 2 6 , 3 3 6 最尤選択部
 3 3 0 D M F C C 算出部
 3 3 2 D M F C C 通常発声認識処理部
 3 3 4 D M F C C 言直し発話認識処理部
 3 5 0 雑音適応化処理部
 3 5 6 M F C C 女声通常発声デコーダ部
 3 5 8 M F C C 男声通常発声デコーダ部
 3 7 0 雑音適応化処理部
 3 7 6 M F C C 女声言直し発話デコーダ部
 3 7 8 M F C C 男声言直し発話デコーダ部
 4 5 2 仮説更新部
 4 5 4 対象単語検索部
 4 5 6 G W P P 算出部
 4 5 8 ウェイト記憶部
 4 6 0 言語モデル記憶部
 4 6 2 単語ラティス作成部
 4 6 4 単語ラティス記憶部
 4 6 6 最高スコア経路探索部
 4 8 0 単語ラティス

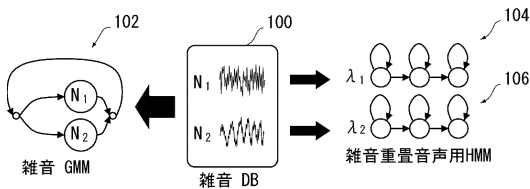
10

20

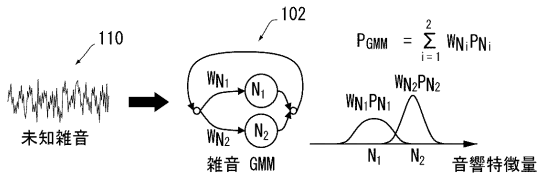
【図 1】



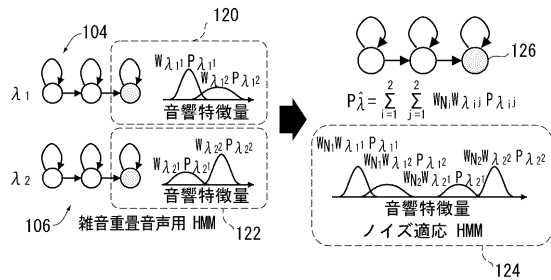
【図 2】



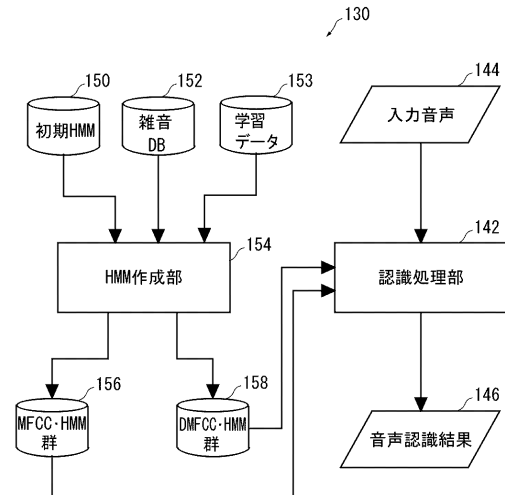
【図 3】



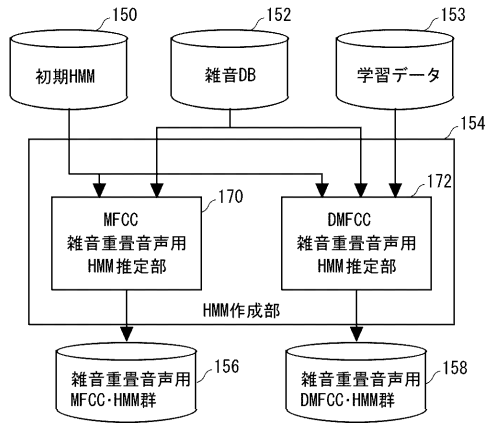
【図 4】



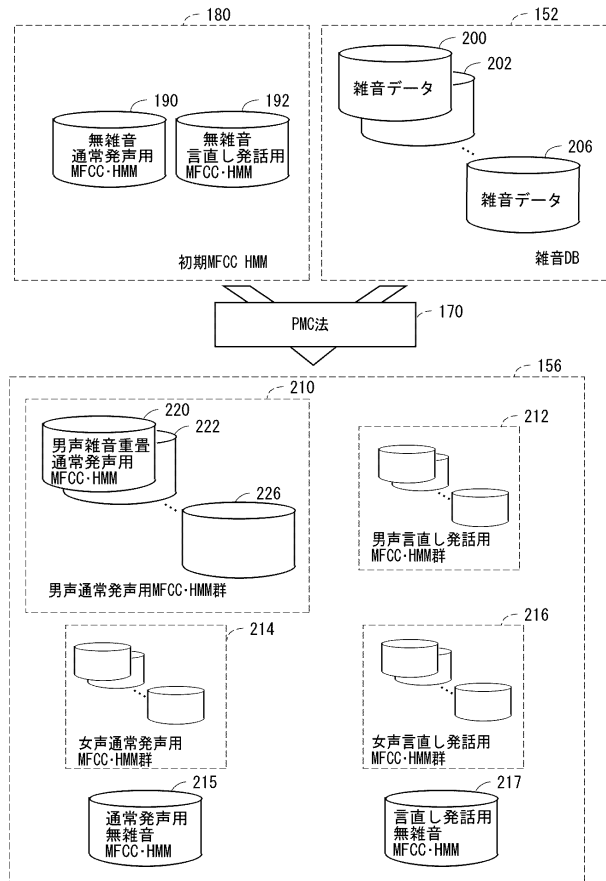
【図 5】



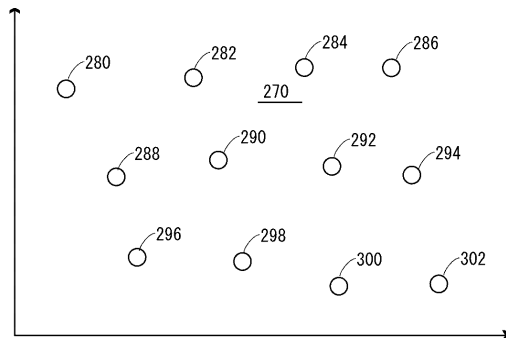
【図 6】



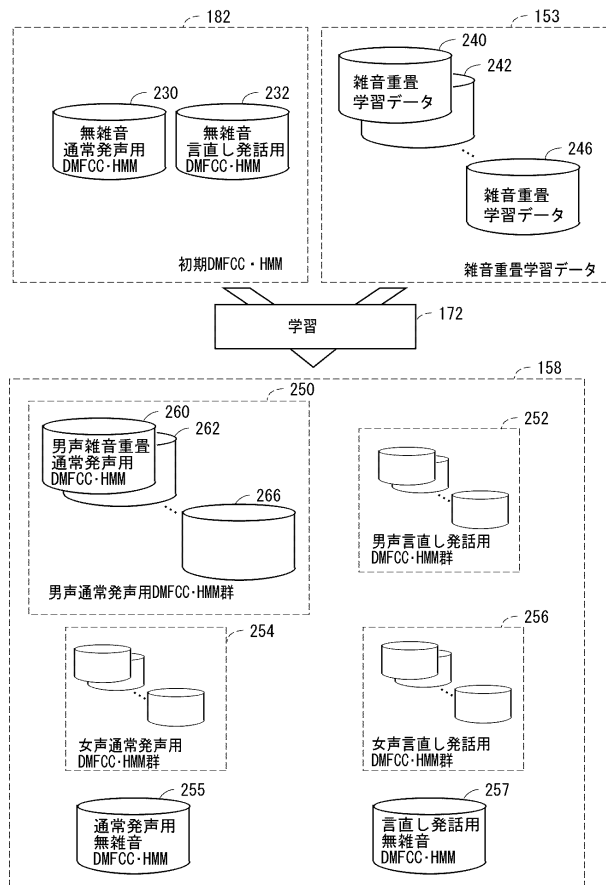
【図 7】



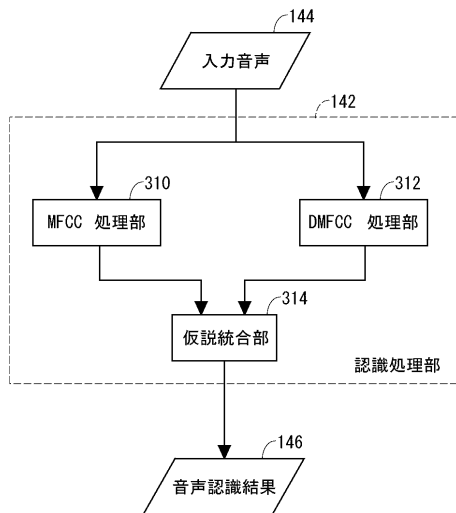
【図 8】



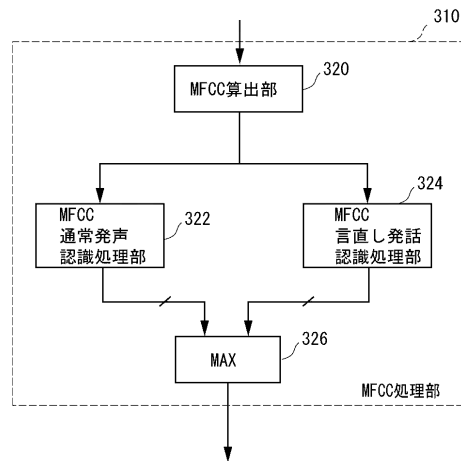
【図 9】



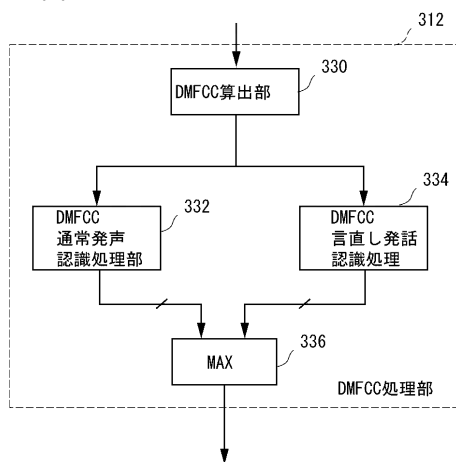
【図 10】



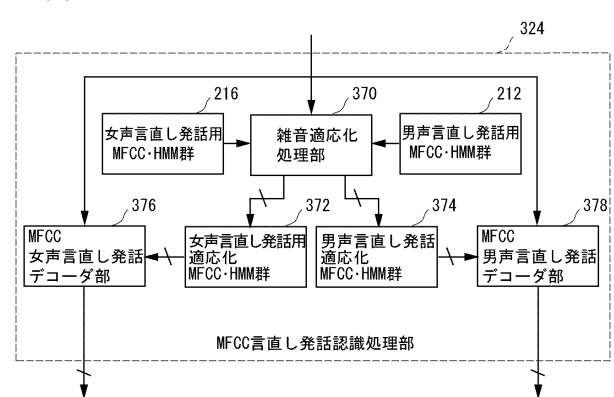
【図 11】



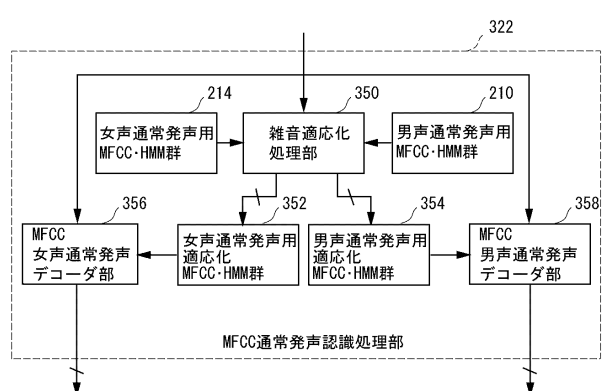
【図 12】



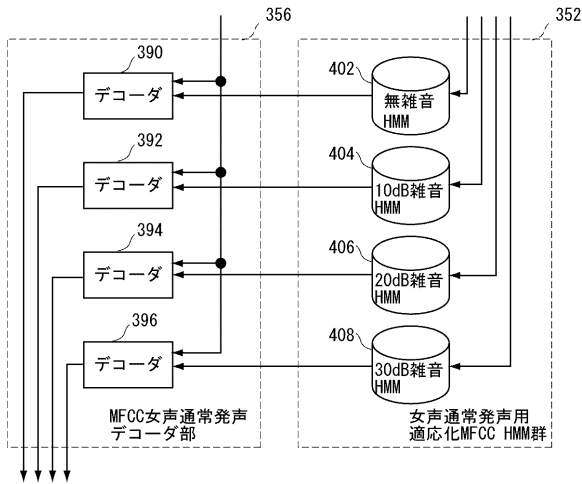
【図 14】



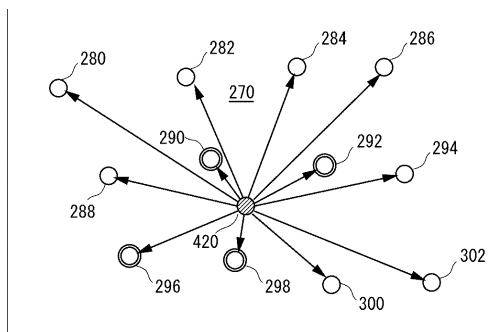
【図 13】



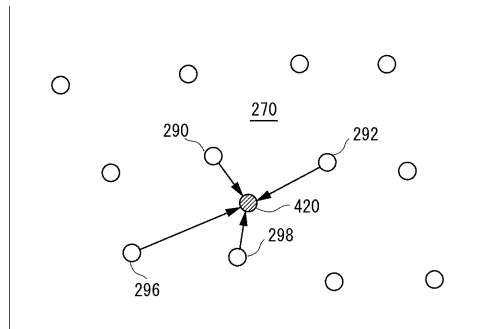
【図 15】



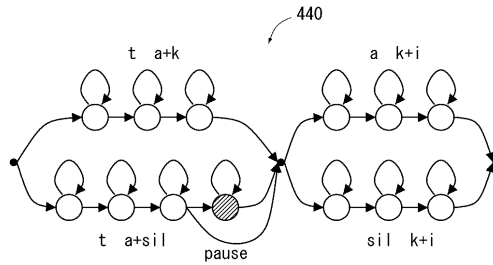
【図 16】



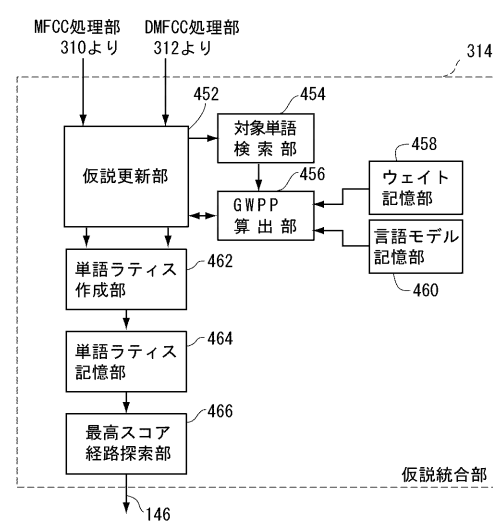
【図 17】



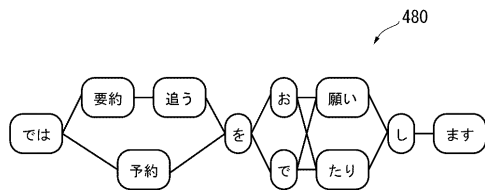
【図 18】



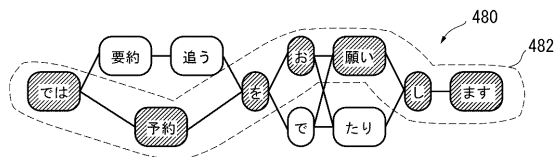
【図 19】



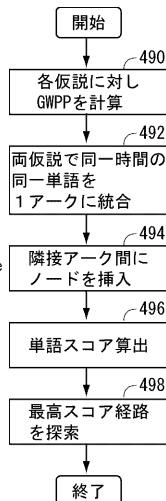
【図 21】



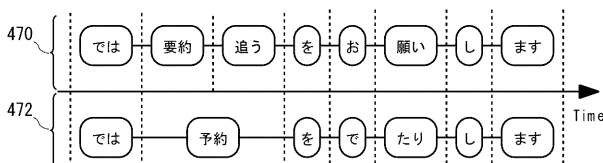
【図 22】



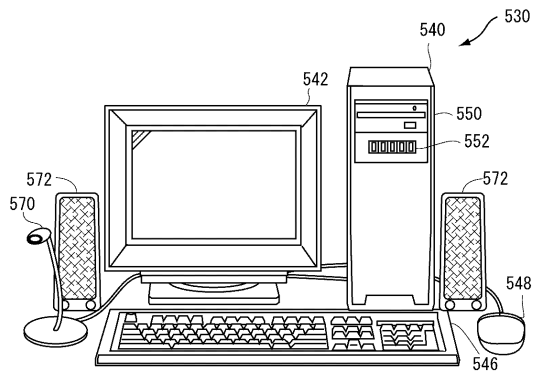
【図 23】



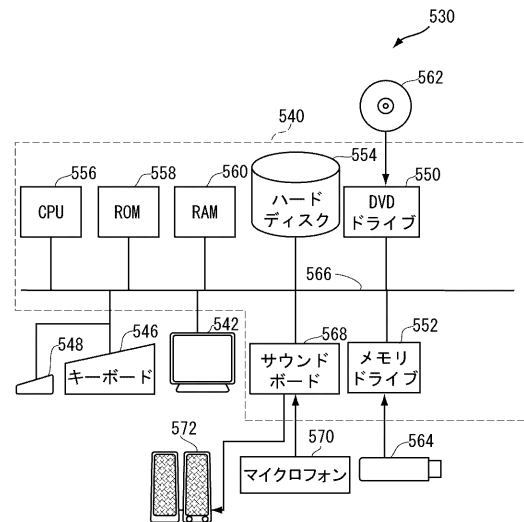
【図 20】



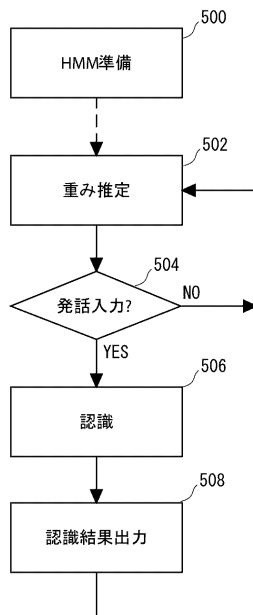
【図 24】



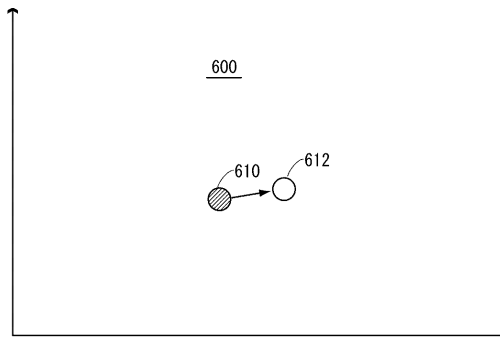
【図 25】



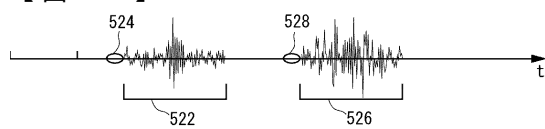
【図 26】



【図 28】



【図 27】



フロントページの続き

(72)発明者 中村 哲

京都府相楽郡精華町光台二丁目2番地2 株式会社国際電気通信基礎技術研究所内

審査官 毛利 太郎

(56)参考文献 特開2005-221678(JP,A)

特開2005-164837(JP,A)

特開2003-345388(JP,A)

松田 繁樹 Shigeki Matsuda, 雑音や発話スタイルの変動に頑健な日本語大語彙連続音声認識

LVCSR Robust to Noise and Speaking Styles, 情報処理学会研究報告 Vol. 2004

No. 15 IPSJ SIG Technical Reports, 日本, 社団法人情報処理学会 Information Processing Society of Japan, 2004年 2月 7日, p.37-44

Wai-Kit Lo, Frank K. SOONG, 中村哲, 一般化事後確率を用いた異なるレベルの大語彙連続音声認識出力の検証, 情報処理学会 研究報告 2004-SLP-54, 日本, 社団法人情報処理学会, 2004年12月22日, pp.259-264

中川聖一, 3.3.4 継続時間の分布を持つHMM, 確率モデルによる音声認識, 日本, 電子情報通信学会, 1988年 7月 1日, p74, p100

山本 博史 Hirofumi YAMAMOTO, 単語適合率最大基準に基づく複数システムの統合 Combination of Multiple Recognizer Outputs Based on Maximum Word Acceptance Rate Path Selection, 電子情報通信学会技術研究報告 Vol. 102 No. 160 IEICE Technical Report, 日本, 社団法人電子情報通信学会 The Institute of Electronics, Information and Communication Engineers, 2002年 6月21日, p.43-47

(58)調査した分野(Int.Cl., DB名)

G10L 15/00 - 17/00