

(19) 日本国特許庁(JP)

(12) 特 許 公 報(B2)

(11) 特許番号

特許第6246777号
(P6246777)

(45) 発行日 平成29年12月13日(2017.12.13)

(24) 登録日 平成29年11月24日(2017.11.24)

(51) Int.Cl.	F I
G 1 0 L 13/10 (2013.01)	G 1 0 L 13/10 1 1 2 B
G 1 0 L 13/033 (2013.01)	G 1 0 L 13/033 1 0 2 A
G 1 0 L 13/06 (2013.01)	G 1 0 L 13/033 1 0 2 Z
	G 1 0 L 13/06 1 3 0
	G 1 0 L 13/10 1 1 2 C
請求項の数 15 (全 30 頁) 最終頁に続く	

(21) 出願番号	特願2015-228796 (P2015-228796)	(73) 特許権者	000003078
(22) 出願日	平成27年11月24日(2015.11.24)		株式会社東芝
(62) 分割の表示	特願2014-241271 (P2014-241271) の分割		東京都港区芝浦一丁目1番1号
原出願日	平成25年3月14日(2013.3.14)	(74) 代理人	100108855
(65) 公開番号	特開2016-66088 (P2016-66088A)		弁理士 蔵田 昌俊
(43) 公開日	平成28年4月28日(2016.4.28)	(74) 代理人	100103034
審査請求日	平成28年2月17日(2016.2.17)		弁理士 野河 信久
(31) 優先権主張番号	1204502.7	(74) 代理人	100075672
(32) 優先日	平成24年3月14日(2012.3.14)		弁理士 峰 隆司
(33) 優先権主張国	英国 (GB)	(74) 代理人	100153051
			弁理士 河野 直樹
		(74) 代理人	100140176
			弁理士 砂川 克
		(74) 代理人	100179062
			弁理士 井上 正
		最終頁に続く	

(54) 【発明の名称】 音声合成方法、装置及びプログラム

(57) 【特許請求の範囲】

【請求項 1】

テキストを入力し、

入力された前記テキストから音声合成した音声に対し感情を付与するための感情特性を複数選択し、

選択された複数の前記感情特性に基づいて前記音声を出力し、

入力された前記テキストを音響単位のシーケンスに分割することと、

音響モデルを使用して、前記音響単位のシーケンスを音声ベクトルのシーケンスに変換することと、を含み、

前記モデルは、音響単位を音声ベクトルに関連付ける確率分布を記述する複数のモデル・パラメータを有し、

選択された音声特性における各々の確率分布の所定のタイプのパラメータは、同一のタイプのパラメータの加重和として表現され、使用される加重和は、音声特性依存であり、

前記音響単位のシーケンスを音声ベクトルのシーケンスに変換することは、選択された前記音声特性のための音声特性依存加重値を検索することを含み、

前記パラメータは、各クラスターにおいて提供され、

前記音声特性依存加重値は、各クラスターごとに検索される、音声合成方法。

【請求項 2】

各々のクラスターは、少なくとも一つの決定木を含み、

前記決定木は、言語上の相違、音声上の相違又は韻律上の相違のうちの少なくとも一つ

10

20

に関連する質問に基づいている、請求項 1 に従う音声合成方法。

【請求項 3】

前記クラスターの前記決定木の間に、構造における相違が存在する、請求項 2 に従う音声合成方法。

【請求項 4】

前記複数の音声特性は、異なる話者音声、異なる話者スタイル、異なる話者感情又は異なるアクセントのうちの少なくとも一つから選択される、請求項 1 に従う音声合成方法。

【請求項 5】

前記確率分布は、ガウス分布、ポアソン分布、ガンマ分布、スチューデント t 分布又はラプラス分布から選択される、請求項 1 に従う音声合成方法。

10

【請求項 6】

音声特性を選択することは、入力を提供することを含み、該入力は、前記音声特性依存加重値が該入力を介して選択されることを可能にする、請求項 1 に従う音声合成方法。

【請求項 7】

音声特性を選択することは、出力される前記テキストから、使用されるべき前記音声特性依存加重値を予測することを含む、請求項 1 に従う音声合成方法。

【請求項 8】

音声特性を選択することは、話者のタイプに関する外部情報から、使用されるべき前記音声特性依存加重値を予測することを含む、請求項 1 に従う音声合成方法。

【請求項 9】

20

音声特性を選択することは、音声を含んでいる音声入力を受信することと、前記音声入力の前記音声の前記音声特性をシミュレートするために前記音声特性依存加重値を変更することを含む、請求項 1 に従う音声合成方法。

【請求項 10】

音声特性を選択することは、複数の予め記憶された複数の加重値セットから、ランダムに一つの加重値セットを選択することを含み、

それぞれの加重値セットは、すべてのクラスターのための複数の前記音声特性依存加重値を含む、請求項 1 に従う音声合成方法。

【請求項 11】

音声特性を選択することは、

30

入力を受信することと、ここで、前記入力は、複数の値を含む、

前記複数の値を、複数の前記音声特性依存加重値にマッピングすることを含む、請求項 1 に従う音声合成方法。

【請求項 12】

前記値は n 次元の値空間を占有し、前記音声特性依存加重値は w 次元加重値空間を占有し、ここで、 n と w は整数であり、 w は n より大きく、前記変換は前記入力値をより高い次元の空間に変換する、請求項 11 に従う音声合成方法。

【請求項 13】

前記複数の値は、認識できる話者特徴を直接表現する、請求項 12 に従う音声合成方法。

40

【請求項 14】

テキストを入力する入力手段と、

入力された前記テキストから音声合成した音声に対し感情を付与するための感情特性を複数選択する選択手段と、

選択された複数の前記感情特性に基づいて前記音声を出力する出力手段と、

入力された前記テキストを音響単位 of シーケンスに分割する分割手段と、

音響モデルを使用して、前記音響単位 of シーケンスを音声ベクトルのシーケンスに変換する変換手段と、を具備し、

前記モデルは、音響単位を音声ベクトルに関連付ける確率分布を記述する複数のモデル・パラメータを有し、

50

選択された音声特性における各々の確率分布の所定のタイプのパラメータは、同一のタイプのパラメータの加重和として表現され、使用される加重和は、音声特性依存であり、前記変換手段は、選択された前記音声特性のための音声特性依存加重値を検索する検索手段を含み、

前記パラメータは、各クラスターにおいて提供され、

前記音声特性依存加重値は、各クラスターごとに検索される、音声合成装置。

【請求項 15】

コンピュータを、

テキストを入力する入力手段と、

入力された前記テキストから音声合成した音声に対し感情を付与するための感情特性を複数選択する選択手段と、

選択された複数の前記感情特性に基づいて前記音声を出力する出力手段と、

入力された前記テキストを音響単位のシーケンスに分割する分割手段と、

音響モデルを使用して、前記音響単位のシーケンスを音声ベクトルのシーケンスに変換する変換手段として機能させるための音声合成プログラムであって、

前記モデルは、音響単位を音声ベクトルに関連付ける確率分布を記述する複数のモデル・パラメータを有し、

選択された音声特性における各々の確率分布の所定のタイプのパラメータは、同一のタイプのパラメータの加重和として表現され、使用される加重和は、音声特性依存であり、

前記変換手段は、選択された前記音声特性のための音声特性依存加重値を検索する検索手段を含み、

前記パラメータは、各クラスターにおいて提供され、

前記音声特性依存加重値は、各クラスターごとに検索される、音声合成プログラム。

【発明の詳細な説明】

【技術分野】

【0001】

(関連出願への相互参照)

この出願は、2012年3月14日付け提出の英国特許出願第1204502.7号に基づくものであり、また、その優先権の利益を主張する。そして、その内容の全体が参照によって本明細書に組み込まれる。

【0002】

(技術分野)

本明細書で一般に説明される実施形態は、テキスト音声合成システム及び方法に関する。

【背景技術】

【0003】

テキスト音声合成システムは、テキストファイルの受理に応じてオーディオ音声又はオーディオ音声ファイルが出力されるシステムである。

【0004】

テキスト音声合成システムは、多種多様のアプリケーション(例えば、電子ゲーム、電子ブック・リーダー、電子メール・リーダー、衛星ナビゲーション、自動電話システム、自動警報システム)で使用される。より人間らしい声のようにシステムに音を出させる要求が継続して存在する。

【図面の簡単な説明】

【0005】

これから添付の図面を参照して、限定されない実施形態に従うシステム及び方法が説明される。添付の図面において各図は次の通りである。

【図1】図1は、テキスト音声合成システムの概略図である。

【図2】図2は、既知の音声処理システムにより実行されるステップを示すフローチャートである。

10

20

30

40

50

【図 3】図 3 は、ガウス確率関数の概略図である。

【図 4】図 4 は、一実施形態に従った音声処理方法のフローチャートである。

【図 5】図 5 は、音声特性がどのようにして選択され得るかについて示すシステムの概略図である。

【図 6】図 6 は、図 5 のシステムに関するバリエーションである。

【図 7】図 7 は、図 5 のシステムに関する更なるバリエーションである。

【図 8】図 8 は、図 5 のシステムに関する更に他のバリエーションである。

【図 9 a】図 9 a は、更なる実施形態に従った音声処理方法のフローチャートである。

【図 9 b】図 9 b は、図 9 a を参照して説明されるステップの一部の図的表現である。

【図 10】図 10 は、訓練可能なテキスト音声合成システムの概略図である。

【図 11】図 11 は、実施形態に従って音声処理システムを訓練する方法を示すフローチャートである。

【図 12】図 12 は、実施形態により用いられる決定木の概略図である。

【図 13】図 13 は、実施形態に従ったシステムの適応を示すフローチャートである。

【図 14】図 14 は、更なる実施形態に従ったシステムの適応を示すフローチャートである。

【詳細な説明】

【0006】

一つの実施形態において、複数の異なる音声特性 (voice characteristics) をシミュレートするために使用するテキスト音声合成方法が提供される。該方法は、テキストを入力することと、入力された該テキストを音響単位のシーケンスに分割することと、入力された該テキストのために音声特性を選択することと、音響モデルを使用して、該音響単位のシーケンスを音声ベクトル (speech vectors) のシーケンスに変換することと、ここで、該モデルは、音響単位を音声ベクトルに関連付ける確率分布を記述する複数のモデル・パラメータを有する、該音声ベクトルのシーケンスを、選択された該音声特性をもつ音声 (audio) として出力することを含み、選択された該音声特性における各々の確率分布の所定のタイプのパラメータは、同一のタイプのパラメータの加重和として表現され、使用される加重和は、音声特性依存であり、該音響単位のシーケンスを音声ベクトルのシーケンスに変換することは、選択された該音声特性のための該音声特性依存加重値を検索することを含み、該パラメータは、各クラスターにおいて提供され、各々のクラスターは、少なくとも一つのサブクラスターを含み、該音声特性依存加重値は、各クラスターごとに検索され、サブクラスターあたりに一つの加重値が存在する。

【0007】

各々のサブクラスターは、少なくとも一つの決定木を含んでも良く、該決定木は、言語上の相違、音声上の相違又は韻律上の相違のうちの少なくとも一つに関連する質問に基づいている。上記クラスターの決定木の間及び上記サブクラスターの木の間に、構造における相違が存在しても良い。

【0008】

上記確率分布は、ガウス分布、ポアソン分布、ガンマ分布、スチューデント t 分布又はラプラス分布から選択されても良い。

【0009】

一つの実施形態において、上記複数の音声特性は、異なる話者音声 (speaker voices)、異なる話者スタイル (speaker styles)、異なる話者感情 (speaker emotions) 又は異なるアクセント (accents) のうちの少なくとも一つから選択される。音声特性を選択することは、入力を提供することを含んでも良く、該入力は、上記加重値が該入力を介して選択されることを可能にする。また、音声特性を選択することは、出力される上記テキストから、使用されるべき上記加重値を予測することを含んでも良い。また、更なる実施形態において、音声特性を選択することは、話者のタイプに関する外部情報から、使用されるべき上記加重値を予測することを含んでも良い。

【0010】

該方法が新たな音声特性に適應することもまた可能である。例えば、音声特性を選択することは、音声（voice）を含んでいる音声入力（audio input）を受信することと、上記音声入力の上記音声の上記音声特性をシミュレートするために上記加重値を変更することを含んでも良い。

【 0 0 1 1 】

更なる実施形態において、音声特性を選択することは、複数の予め記憶された複数の加重値セットから、ランダムに一つの加重値セットを選択することを含み、それぞれの加重値セットは、すべてのサブクラスターのための複数の上記加重値を含む。

【 0 0 1 2 】

更なる実施形態において、音声特性を選択することは、入力を受信することと、ここで、上記入力は、複数の値を含む、上記複数の値を、複数の上記加重値にマッピングすることを含む。例えば、上記の値は n 次元の値空間を占有し、上記加重値は w 次元加重値空間を占有し、ここで、 n と w は整数であり、 w は n より大きく、上記変換は上記入力値をより高い次元の空間に変換し得る。それら値は、認識できる話者特性（例えば、うれしい声（happy voice）、気が立っている声（nervous voice）、怒った声（angry voice）など）を直接表現しても良い。そして、それら値の空間は、ユーザが又はテキストのコンテキストに関する何らかの他のインジケーションが、出力される声が感情空間の上のどこにあるべきかを示す「感情空間」と考えることができる。これは「感情空間」より非常に大きなディメンションをしばしば有するであろう加重値空間の上へマッピングされる。

【 0 0 1 3 】

他の実施形態において、テキスト音声合成システムをオーディオ・ファイルに含まれる音声特性に適應する方法が提供され、該テキスト音声合成システムは、テキストを入力し、入力された該テキストを音響単位 of シーケンスに分割し、入力された該テキストのために音声特性を選択し、音響モデルを使用して、該音響単位 of シーケンスを音声ベクトルのシーケンスに変換し、ここで、該モデルは、音響単位を音声ベクトルに関連付ける確率分布を記述する複数のモデル・パラメータを有する、該音声ベクトルのシーケンスを、選択された該音声特性をもつ音声として出力するように構成されたプロセッサを含み、選択された該音声特性における各々の確率分布の所定のタイプのパラメータは、同一のタイプのパラメータの加重和として表現され、使用される加重和は、音声特性依存であり、該音響単位 of シーケンスを音声ベクトルのシーケンスに変換することは、選択された該音声特性のための該音声特性依存加重値を検索することを含み、該パラメータは、各クラスターにおいて提供され、各々のクラスターは、少なくとも一つのサブクラスターを含み、該音声特性依存加重値は、各クラスターごとに検索され、サブクラスターあたりに一つの加重値が存在し、該方法は、新たな入力オーディオ・ファイルを受信することと、生成された該音声と該新たなオーディオ・ファイルとの間の類似を最大にするために、該クラスターに適用される該加重値を計算することを含む。

【 0 0 1 4 】

更なる実施形態において、上記新たなオーディオ・ファイルからのデータを使用して新たなクラスターが生成され、生成された上記音声と上記新たなオーディオ・ファイルとの間の上記類似を最大にするために、上記新たなクラスターを含む上記クラスターに適用される上記加重値が計算される。

【 0 0 1 5 】

より厳密に新たなオーディオ・ファイルの音声にマッチするように、上記加重値に加えて話者変換（例えば、CMLLR変換）が適用されても良い。生成された上記音声と上記新たなオーディオ・ファイルとの間の上記類似を最大にするために、上記の線形変換が適用されても良い。新たな話者クラスターを生成することなく適合が行われる場合及び新たな話者クラスターが生成される場合の両方において、追加の変換（extra transform）を適用するこの技術が使用されても良い。

【 0 0 1 6 】

更なる実施形態において、複数の異なる音声特性をシミュレートするために使用される

10

20

30

40

50

テキスト音声合成システムが提供され、該システムは、入力されたテキストを受信するためのテキスト入力と、プロセッサとを含み、該プロセッサは、入力された該テキストを音響単位のシーケンスに分割し、入力された該テキストのための音声特性の選択を可能にし、音響モデルを使用して、該音響単位のシーケンスを音声ベクトルのシーケンスに変換し、ここで、該モデルは、音響単位を音声ベクトルに関連付ける確率分布を記述する複数のモデル・パラメータを有する、該音声ベクトルのシーケンスを、選択された該音声特性をもつ音声として出力するように構成され、選択された該音声特性における各々の確率分布の所定のタイプのパラメータは、同一のタイプのパラメータの加重和として表現され、使用される加重和は、音声特性依存であり、該音響単位のシーケンスを音声ベクトルのシーケンスに変換することは、選択された該音声特性のための該音声特性依存加重値を検索することを含み、該パラメータは、各クラスターにおいて提供され、各々のクラスターは、少なくとも一つのサブクラスターを含み、該音声特性依存加重値は、各クラスターごとに検索され、サブクラスターあたりに一つの加重値が存在する、システム。

10

【 0 0 1 7 】

更なる実施形態において、オーディオ・ファイルにおいて提供される音声特性をもつ音声出力するように構成された、適応性のあるテキスト音声合成システムが提供され、該テキスト音声合成システムは、入力されたテキストを受信し、入力された該テキストを音響単位のシーケンスに分割し、入力された該テキストのために音声特性を選択し、音響モデルを使用して、該音響単位のシーケンスを音声ベクトルのシーケンスに変換し、ここで、該モデルは、音響単位を音声ベクトルに関連付ける確率分布を記述する複数のモデル・パラメータを有する、該音声ベクトルのシーケンスを、選択された該音声特性をもつ音声として出力するように構成されたプロセッサを含み、選択された該音声特性における各々の確率分布の所定のタイプのパラメータは、同一のタイプのパラメータの加重和として表現され、使用される加重和は、音声特性依存であり、該音響単位のシーケンスを音声ベクトルのシーケンスに変換することは、選択された該音声特性のための該音声特性依存加重値を検索することを含み、該パラメータは、各クラスターにおいて提供され、各々のクラスターは、少なくとも一つのサブクラスターを含み、該音声特性依存加重値は、各クラスターごとに検索され、サブクラスターあたりに一つの加重値が存在し、該システムは、クラスター及びサブクラスターにおいて提供される該パラメータと、該サブクラスターのための該加重値とを記憶するように構成されたメモリを更に含み、該システムは、新たな入力オーディオ・ファイルを受信するように更に構成され、該プロセッサは、生成された該音声と該新たなオーディオ・ファイルとの間の類似を最大にするために、該サブクラスターに適用される該加重値を再計算するように構成される。

20

30

【 0 0 1 8 】

本発明の実施形態に従う方法は、ハードウェアでも汎用コンピュータ中のソフトウェアでも実施することができる。本実施形態に従う更なる方法は、ハードとソフトとの組み合わせで実施することができる。本発明の実施形態に従う方法はまた、単一の処理装置、又は複数の処理装置からなる分散ネットワークにより実施することができる。

【 0 0 1 9 】

実施形態に従う幾つかの方法はソフトウェアにより実施することができるので、幾つかの実施形態は任意の適したキャリア媒体上の汎用コンピュータに提供されるコンピュータコードを含む。キャリア媒体は、例えばフロッピー（登録商標）ディスク、CD ROM、磁気デバイス若しくはプログラマブル・メモリ・デバイスのような任意の記憶媒体、又は、例えば任意の信号（例えば、電気的信号、光学的信号若しくはマイクロ波信号）のような任意の一時的な媒体を含むことができる。

40

【 0 0 2 0 】

図 1 は、テキスト音声合成システム 1 を示す。テキスト音声合成システム 1 は、プログラム 5 を実行するプロセッサ 3 を含む。テキスト音声合成システム 1 は、記憶装置 7 を更に含む。記憶装置 7 は、テキストを音声に変換するプログラム 5 により使用されるデータを記憶する。テキスト音声合成システム 1 は、入力モジュール 11 及び出力モジュール 1

50

3を更に含む。入力モジュール11は、テキスト入力15に接続される。テキスト入力15は、テキストを受ける。テキスト入力15は、例えば、キーボードであっても良い。あるいは、テキスト入力15は、外部記憶媒体又はネットワークから、テキストデータを受信するための手段であっても良い。

【0021】

出力モジュール13に接続されるのは、音声用出力17である。音声出力17は、テキスト入力15へ入力されるテキストから変換された音声信号を出力するために使用される。音声出力17は、例えば、直接の音声出力（例えば、スピーカ）であっても良いし、又は、記憶媒体、ネットワークなどに送信され得るオーディオ・データ・ファイル用の出力であっても良い。

10

【0022】

使用するときは、テキスト音声合成システム1は、テキスト入力15を通してテキストを受け取る。プロセッサ3上で実行されるプログラム5は、記憶装置7に記憶されたデータを使用して、テキストを音声データに変換する。音声は、出力モジュール13を介して音声出力17へ出力される。

【0023】

これから図2を参照して単純化したプロセスが説明される。最初のステップS101において、テキストが入力される。テキストは、キーボード、タッチ・スクリーン、テキスト予測機能又は同様のものを介して入力されても良い。その後、テキストは、音響単位のシーケンスに変換される。これらの音響単位は、音素又は書記素であっても良い。該単位は、コンテキスト依存（例えば、選択された音素に加えて先行する音素及び後続する音素も考慮に入れるトライフォン）であっても良い。該テキストは、当該技術において周知の（本明細書では更に説明されない）技術を使用して、音響単位のシーケンスに変換される。

20

【0024】

S105において、音響単位を音声パラメータに関連付ける確率分布が検索される。この実施形態において、確率分布は、平均及び分散により定義されるガウス分布であることがある。例えばポアソン分布、学生t分布、ラプラス分布又はガンマ分布のような他の分布を使用することが可能であるが、それらのうちの幾つかは、平均及び分散とは異なる変数により定義される。

30

【0025】

各々の音響単位が、音声ベクトル又は当該技術の専門用語を使用する「観測（observation）」に対して明確な一対一の対応を有することはとても有り得ない。多くの音響単位は、類似する方法で発音され、また、周囲の音響単位によって、或いは、単語若しくは文におけるそれらの位置によって、影響を受け、又は、異なる話者により異なった風に発音される。したがって、各々の音響単位は、音声ベクトルに関連付けられる確率を有するのみであり、また、テキスト音声合成システムは、多くの確率を計算して、音響単位のシーケンスを与えられた複数の観測のうち、最も起こり得るシーケンスを選択する。

【0026】

ガウス分布は図3に示される。図3は、音声ベクトルに関係する音響単位の確率分布であるものとして考えることができる。例えば、Xとして示された音声ベクトルは、図3に示される分布を有する音素又は他の音響単位に対応する確率P1を有する。

40

【0027】

ガウス分布の形状及び位置は、その平均及び分散により定義される。これらのパラメータは、システムの訓練の間に決定される。

【0028】

その後、ステップS107において、これらのパラメータが音響モデルにおいて使用される。この説明において、音響モデルは、隠れマルコフモデル（HMM）である。しかしながら、他のモデルを使用することもできる。

【0029】

50

音声システムのテキストは、音響単位（すなわち、音素、書記素、単語又はその部分）を音声パラメータに関連付ける多数の確率密度関数を記憶する。ガウス分布が一般に使用されるように、これらは一般にガウシアン又はコンポーネントと呼ばれる。

【 0 0 3 0 】

隠れマルコフモデル又は他のタイプの音響モデルにおいて、特定の音響単位に関するすべての可能性のある音声ベクトルの確率が考慮される必要がある。そして、その音響単位のシーケンスに最大の可能性をもって対応する音声ベクトルのシーケンスが考慮される。これは、二つの単位が互いに影響を及ぼす方法（way）を考慮に入れる、シーケンスのすべての音響単位にわたる、大域的最適化（global optimization）を意味する。その結果、複数の音響単位からなるシーケンスが考慮される場合に、特定の音響単位に対する最有望な音声ベクトルが最良の音声ベクトルにならないことがあり得る。

10

【 0 0 3 1 】

音声ベクトルのシーケンスが決定されると、ステップ S 1 0 9 において、音声が出力される。

【 0 0 3 2 】

図 4 は、一実施形態に従ったテキスト音声合成システムのためのプロセスのフローチャートである。ステップ S 2 0 1 において、図 2 を参照して説明された方法と同じ方法で、テキストが受信される。その後、ステップ S 2 0 3 において、該テキストは、音素、書記素、コンテキスト依存の音素又は書記素、及び、単語又はその部分であり得る音響単位のシーケンスに変換される。

20

【 0 0 3 3 】

図 4 のシステムは、幾つかの異なる音声特性を使用して音声を出力することができる。例えば、実施形態において、特性は、うれしい（happy）、悲しい（sad）、怒った（angry）、気が立っている（nervous）、穏やかな（calm）、威圧する（commanding）などのように聞こえる声から選択されても良い。

【 0 0 3 4 】

ステップ S 2 0 5 において、要求される音声特性が判定される。これは、幾つかの異なる方法によってなされても良い。選択された音声特性を判定する幾つかの可能な方法の例が、図 5 ～ 8 を参照して説明される。

【 0 0 3 5 】

30

図 4 を参照して説明される方法において、各々のガウシアン・コンポーネントが平均及び分散により記述される。幾つかの実施形態においては、それぞれであろう複数の異なる状態が、ガウス分布を用いてモデル化されるであろう。例えば、一つの実施形態では、テキスト音声合成システムは、複数のストリームを含む。それらのようなストリームは、一つ又は複数のスペクトル・パラメータ（Spectrum）、基本周波数の対数（ $\text{Log } F_0$ ）、 $\text{Log } F_0$ の一次微分（ $\text{Delta Log } F_0$ ）、 $\text{Log } F_0$ の二次微分（ $\text{Delta-Delta Log } F_0$ ）、バンド非周期性パラメータ（Band aperiodicity parameters）（BAP）、継続期間（duration）などから選択されても良い。ストリームはまた、クラス（例えば、無音（silence）（sil）、短いポーズ（short pause）（pau）及び音声（speech）（spe）など）に更に分けられても良い。一つの実施形態では、ストリーム及びクラスのそれぞれからのデータは、HMM を使用してモデル化される。HMM は、異なる数の状態を含んでも良い。例えば、一つの実施形態において、上記のストリーム及びクラスのうちの一部分からのデータをモデル化するために、5 状態 HMM（5 state HMMs）が用いられても良い。ガウシアン・コンポーネントは、各 HMM 状態ごとに決定される。

40

【 0 0 3 6 】

図 4 のシステムにおいて、選択された音声特性をもつガウス分布の平均は、ガウス分布の非依存平均（independent means）の加重和として表現される。したがって、次のようになる。

【数 1】

$$\mu_m^{(s)} = \sum_i \lambda_i^{(s)} \mu_{c(m,i)} \quad \dots \text{式 (1)}$$

【0037】

ここで、 $\mu_m^{(s)}$ は、選択された話者音声 s におけるコンポーネント m のための平均であり、 $i \in \{1, \dots, P\}$ は、クラスターのインデックスであり、 P は、クラスターの総数であり、 $\lambda_i^{(s)}$ は、話者 s のための第 i 番目のクラスターの話者依存補間加重値 (speaker dependent interpolation weight) であり、 $\mu_{c(m,i)}$ は、クラスター i におけるコンポーネント m のための平均である。複数のクラスターのうちの一つ (通常、クラスター $i = 1$) に対して、すべての加重値が常に 1.0 にセットされる。このクラスターは、“バイアス・クラスター” と呼ばれる。それぞれのクラスターは、少なくとも一つの決定木を含む。決定木は、クラスター中の各コンポーネントごとに存在する。表現を単純化するために、 $c(m, i) \in \{1, \dots, N\}$ は、クラスター i のための平均ベクトル決定木におけるコンポーネント m のための総合リーフ・ノード・インデックスを示す。 N は、すべてのクラスターの決定木にわたるリーフ・ノードの総数である。決定木の詳細は、後で説明される。

10

【0038】

ステップ S 207 において、システムは、アクセス可能な方法で記憶される平均及び分散を検索する。

20

【0039】

ステップ S 209 において、システムは、それら平均について音声特性依存加重値を検索する。それら音声特性依存加重値は、それら平均が検索される前又は後に検索されても良いことは当業者により認識されるであろう。

【0040】

したがって、ステップ S 209 の後で、それら音声特性依存平均を得ること、すなわち、それら平均を使用すること及びそれら加重値を適用すること、が可能である。その後、これらは、図 2 中のステップ S 107 を参照して説明された方法と同じ方法でステップ S 211 中の音響モデルにおいて使用される。その後、ステップ S 213 において、音声が出力される。

30

【0041】

音声特性非依存平均は、クラスター化される。一つの実施形態では、それぞれのクラスターは、少なくとも一つの決定木を含み、木において使用される決定は、言語上の変動、音声上の変動又は韻律上の変動に基づく。一つの実施形態では、決定木は、クラスターのメンバーである各コンポーネントごとに存在する。韻律上のコンテキスト、音声上のコンテキスト及び言語上のコンテキストは、最終的な音声波形に影響を及ぼす。音声上のコンテキストは、典型的には、声道に影響を及ぼし、韻律上のコンテキスト (例えば音節) 及び言語上のコンテキスト (例えば単語の品詞) は、例えば継続時間 (リズム) および基本周波数 (トーン) のような韻律に影響を及ぼす。それぞれのクラスターは、1 又は複数のサブクラスターを含んでも良い。それぞれのサブクラスターは、それら決定木のうちの少なくとも一つを含む。

40

【0042】

上記は、各サブクラスターごとに加重値を検索することも又は各クラスターごとに加重値ベクトルを検索することも考慮することができる。ここで、加重値ベクトルの要素は、各サブクラスターのための加重値である。

【0043】

一つの実施形態に従って以下の構成が使用されても良い。このデータをモデル化するために、この実施形態では、5 状態 HMM が使用される。この例に関して、データは、無音、短いポーズ、音声の三つのクラスに分けられる。この特定の実施形態において、サブク

50

ラスターごとの決定木及び加重値の割り当ては、次のとおりである。

【 0 0 4 4 】

この特定の実施形態では、クラスターごとに次のストリームが使用される。

Spectrum : 1つのストリーム、5つの状態、状態ごとに1つの木×3クラス

LogF0 : 3つのストリーム、ストリームごとに5つの状態、状態及びストリームごとに、1つの木×3クラス

BAP : 1つのストリーム、5つの状態、状態ごとに1つの木×3クラス

継続期間 : 1つのストリーム、5つの状態、1つの木×3クラス（各木は、すべての状態にわたって共有される）

合計 : $3 \times 26 = 78$ の決定木

10

上記に関して、次の加重値が、音声特性（例えば話者）ごとに、各々のストリームに適用される。

Spectrum : 1つのストリーム、5つの状態、ストリームごとに1つの加重値×3クラス

LogF0 : 3つのストリーム、ストリームごとに5つの状態、ストリームごとに1つの加重値×3クラス

BAP : 1つのストリーム、5つの状態、ストリームごとに1つの加重値×3クラス

継続時間 : 1つのストリーム、5つの状態、状態及びストリームごとに1つの加重値×3クラス

合計 : $3 \times 10 = 30$ の加重値

20

この例で示されるように、異なる決定木（spectrum）に同一の加重値を割り当てること、あるいは、同一の決定木（継続時間）に2以上の加重値を割り当てること、又は、任意の他の組み合わせが、可能である。本明細書で使用されるように、同一の加重値が適用されるべき決定木は、サブクラスターを形成するために考慮される。

【 0 0 4 5 】

一つの実施形態において、選択された音声特性をもつガウス分布の平均は、複数のガウシアン・コンポーネントの平均の加重和として表現される。ここで、該加重和は、各々のクラスターから1つずつの平均を使用する。該平均は、現在処理されている音響単位の韻律上のコンテキスト、言語上のコンテキスト及び音声上のコンテキストに基づいて選択される。

30

【 0 0 4 6 】

図5は、音声特性を選択する可能な方法を示す。ここでは、ユーザは、例えば、スクリーン上でポイントをドラッグアンドドロップするためのマウス、数量を入力するためのキーボードなどを使用して、加重値を直接選択する。図5において、マウス、キーボード又は同様のものを含む選択ユニット251は、ディスプレイ253を使用して、加重値を選択する。この例では、ディスプレイ253は、加重値を示すレーダー・チャートを含む。ユーザは、レーダー・チャートを介して様々なクラスターの優位性（dominance）を変えるために、選択ユニット251を使用することができる。他の表示方法が使用され得ることは、当業者により認識されるであろう。

【 0 0 4 7 】

40

幾つかの実施形態においては、加重値は、それら自身の空間、すなわち「加重値空間」に射影されることができる。加重値空間は、初期的にそれぞれのディメンションを表す加重値をもつ。この空間は、異なる空間に再配置されることができる。異なる空間のディメンションは、異なる音声属性（voice attributes）を表す。例えば、モデル化された音声特性が表現（expression）であるならば、一つのディメンションは、うれしい音声特性（happy voice characteristics）を示し、他のディメンションは、気が立っている（nervous）音声特性を示すなどとしても良い。ユーザは、うれしい声のディメンション（happy voice dimension）に関して、この音声特性が優位を占めるように、加重値を増加させるように選択しても良い。この場合、新たな空間のディメンションの数は、オリジナルの加重値空間のそれより少ない。

50

【 0 0 4 8 】

そして、オリジナルの空間 $\lambda^{(s)}$ の加重値ベクトルは、新たな空間 $\alpha^{(s)}$ の座標ベクトルの関数として得ることができる。

【 0 0 4 9 】

一つの実施形態では、低減されたディメンション加重値空間の上への、このオリジナルの加重値空間の射影は、タイプ $\lambda^{(s)} = H \alpha^{(s)}$ の一次方程式を使用して形成される。ここで、H は、射影行列である。一つの実施形態では、行列 H は、マニュアルで選択された d の代表的な話者に対するオリジナルの $\lambda^{(s)}$ を、その列にセットするように定義される。ここで、d は、新たな空間の要求されるディメンションである。加重値空間の次元を減らすために、あるいは、 $\alpha^{(s)}$ の値が幾つかの話者についてあらかじめ定義される場合に、制御 α 空間をオリジナルの λ 加重値空間にマッピングする関数を自動的に見出すために、他の技術を使用し得る。

10

【 0 0 5 0 】

更なる実施形態では、システムは、予め定められた複数の加重値ベクトル・セットを記憶するメモリにより提供される。各々のベクトルは、テキストが、異なる音声特性で出力されていることを可能にするようにデザインされても良い。例えば、うれしい声、激怒した声など。図 6 に、そのような実施形態に従うシステムが示される。ここでは、ディスプレイ 253 は、選択ユニット 251 により選択され得る異なる音声属性を示す。

【 0 0 5 1 】

システムは、予め定められた複数セットの属性に基づいて、話者出力の選択肢のセットを示しても良い。ユーザは、それから、求められる話者を選択しても良い。

20

【 0 0 5 2 】

更なる実施形態において、図 7 で示されるように、システムは、自動的に加重値を判定する。例えば、システムは、それが命令又は質問であると認識するテキストに対応する音声出力する必要がある場合がある。システムは、電子ブックを出力するように構成される場合がある。システムは、語り手とは対照的に本の登場人物により何かが話されているときにテキストから、例えば引用符から、認識しても良く、また、新たな音声特性を出力に導入するために、加重値を変更しても良い。同様に、システムは、テキストが繰り返されるかどうか認識するように構成されても良い。そのような状況において、音声特性は、2 番目の出力を変更しても良い。更に、システムは、テキストが、うれしい瞬間又は不安な瞬間に言及するかどうか、そして、テキストが、適切な音声特性で出力されるかどうか、認識するように構成されても良い。

30

【 0 0 5 3 】

上記のシステムにおいて、テキストにおいてチェックされるべき属性及びルールを記憶するメモリ 261 が、提供される。入力テキストは、ユニット 263 によりメモリ 261 に提供される。テキストに対してルールがチェックされ、そして、音声特性のタイプに関する情報が選択ユニット 265 に渡される。選択ユニット 265 は、それから、選択された音声特性のための加重値を検索する。

【 0 0 5 4 】

上記のシステム及び考慮点はまた、ゲーム中のキャラクターが話すコンピュータ・ゲームにおいて使用されるシステムに適用されても良い。

40

【 0 0 5 5 】

更なる実施形態において、システムは、更なるソースから出力されるテキストに関する情報を受信する。図 8 に、そのようなシステムの例が示される。例えば、電子ブックの場合に、システムは、テキストの特定の部分がどのようにして出力されるべきかについて示す入力を受信しても良い。

【 0 0 5 6 】

コンピュータ・ゲームにおいて、システムは、話しているキャラクターが、負傷したかどうか、キャラクターが隠れていて、ささやかなければならないかどうか、誰かの注目を集めているキャラクターが、ゲームのステージをうまく完了したかどうか、などを、ゲー

50

ムから判定することができる。

【 0 0 5 7 】

図 8 のシステムにおいて、テキストがどのようにして出力されるべきかという詳しい情報が、ユニット 2 7 1 から受信される。ユニット 2 7 1 は、それから、この情報をメモリ 2 7 3 に送信する。メモリ 2 7 3 は、それから、音声 (voice) がどのように出力されるべきかに関する情報を検索して、これをユニット 2 7 5 に送信する。ユニット 2 7 5 は、それから、要求される音声出力のための加重値を検索する。

【 0 0 5 8 】

上記に加えて、本方法は、M L L R、C M L L R 変換又は同様のものを使用することによって、更に音声変換 (voice transform) を実装する。具体的には、モデル化された音声特性が話者変動性 (speaker variation) であるとき、この追加の変換は、クラスターの加重値により提供される任意の話者変動に加えて、追加のモデリング能力 (extra modeling power) を加える。この追加の変換を使用するプロセスは、図 9 a 及び図 9 b で説明される。

【 0 0 5 9 】

図 9 a では、ステップ S 2 0 6 において、話者音声を選択される。話者音声は、既知の話者変換により実装可能な、複数の予め記憶された話者プロファイルから選択される。選択された話者プロファイルは、システムの初期セットアップの間に判定可能であり、毎回システムが使用されるとは限らない。

【 0 0 6 0 】

前に説明されたように、システムは、それから、S 2 0 7 においてモデル・パラメータを検索し、ステップ S 2 0 9 において要求に応じ話者加重値を検索する。

【 0 0 6 1 】

システムが、要求される話者を知っているとき、ステップ S 2 1 0 において、システムは追加の話者変換を検索することができる。それから、ステップ S 2 1 1 において、話者依存加重値及び話者変換と一緒に適用される。前述のとおり、ステップ S 2 1 2 で音声ベクトルのセットが判定され、そして、ステップ S 2 1 3 で音声出力される。この具体例では、上記変換は、音声ベクトルの生成の前に、モデルに適用される。

【 0 0 6 2 】

図 9 b は、図 9 a に関して説明されるプロセスの概略図である。図 9 a のステップ S 2 0 9 において、話者加重値が検索される。これらの加重値は、図 9 b の決定木 4 0 1 に適用される。各々の決定木からの加重値平均は、4 0 3 において合計される。4 0 5 において、話者変換 (使用されるならば) が適用され、そして、4 0 7 において、最終的な話者モデルが出力される。

【 0 0 6 3 】

次に、図 1 0 及び図 1 1 を参照して、本発明の一実施形態に従ったシステムの訓練が説明される。

【 0 0 6 4 】

図 1 0 のシステムは、図 1 を参照して説明されたシステムに類似している。したがって、不必要な重複を避けるために、同様の特徴を示すために同様の参照番号が使用される。

【 0 0 6 5 】

図 1 を参照して説明された特徴に加えて、図 1 0 は、音声入力 2 3 及び音声入力モジュール 2 1 を更に含む。システムを訓練する場合、テキスト入力 1 5 を介して入力されているテキストにマッチする音声入力を有することが必要である。

【 0 0 6 6 】

隠れマルコフモデル (H M M) に基づく音声処理システムにおいて、H M M はしばしば次のように表現される。

10

20

30

40

【数 2】

$$M = (A, B, \Pi) \quad \dots \text{式 (2)}$$

【0067】

ここで、A は状態遷移確率分布であり、次のようである。

【数 3】

$$A = \{a_{ij}\}_{i,j=1}^N$$

10

【0068】

また、B は状態出力確率分布であり、次のようである。

【数 4】

$$B = \{b_j(\mathbf{o})\}_{j=1}^N$$

【0069】

また、 Π は初期状態確率分布であり、次のようである。

20

【数 5】

$$\Pi = \{\pi_i\}_{i=1}^N$$

【0070】

ここで、N は、HMM における状態の数である。

【0071】

テキスト音声合成システムにおいて HMM がどのように使用されるかについては、当該技術では周知であり、ここでは説明されない。

30

【0072】

現在の実施形態において、状態遷移確率分散 A 及び初期状態確率分布は、当該技術において周知の手続きに従って決定される。したがって、この説明の残りは、状態出力確率分布に関係している。

【0073】

一般に、テキスト音声合成システムにおいて、モデルセット M における第 m 番目のガウシアン・コンポーネントからの状態出力ベクトル又は音声ベクトル $\mathbf{o}(t)$ は、次のようになる。

【数 6】

$$P(\mathbf{o}(t) | m, s, M) = N(\mathbf{o}(t); \boldsymbol{\mu}_m^{(s)}, \boldsymbol{\Sigma}_m^{(s)})$$

40

…式 (3)

【0074】

ここで、 $\boldsymbol{\mu}_m^{(s)}$ と $\boldsymbol{\Sigma}_m^{(s)}$ は、話者 s のための第 m 番目のガウシアン・コンポーネントの平均と共分散である。

【0075】

従来のテキスト音声合成システムを訓練する場合の目標は、与えられた観測シーケンスに対する尤度を最大化するモデル・パラメータ・セット M を推定することである。従来の

50

モデルでは、単一の話者が存在し、したがって、モデル・パラメータ・セットは、すべてのコンポーネント m について、 $\mu^{(s)}_m = \mu_m$ 及び $\Sigma^{(s)}_m = \Sigma_m$ である。

【 0 0 7 6 】

いわゆる最尤 (M L) 基準に純粹に分析的に基づいて上記のモデルセットを得ることは可能でないので、従来、その問題は、バウム・ウェルチ・アルゴリズムと大抵呼ばれる期待値最大化 (E M) アルゴリズムとして知られている反復アプローチを使用することによって対処される。ここで、次のような補助関数 (“ Q ” 関数) が得られる。

【 数 7 】

$$Q(M, M') = \sum_{m,t} \gamma_m(t) \log p(o(t), m | M) \quad \dots \text{式 (4)} \quad 10$$

【 0 0 7 7 】

ここで、 $\gamma_m(t)$ は、観測 $o(t)$ を生成するコンポーネント m の事後確率であり、現在のモデル・パラメータは M' 、 M は新たなパラメータ・セットとする。各々の反復の後で、パラメータ・セット M' は、 $Q(M, M')$ を最大化する新たなパラメータ・セット M と置き換えられる。 $p(o(t), m | M)$ は、例えばGMM、HMMなどのような生成モデルである。

【 0 0 7 8 】

現在の実施形態において、次式の状態出力ベクトルを有するHMMが使用される。

20

【 数 8 】

$$P(o(t) | m, s, M) = N(o(t); \hat{\mu}_m^{(s)}, \hat{\Sigma}_{v(m)}^{(s)}) \quad \dots \text{式 (5)}$$

【 0 0 7 9 】

ここで、 $m \in \{1, \dots, MN\}$ 、 $t \in \{1, \dots, T\}$ 、及び、 $s \in \{1, \dots, S\}$ は、それぞれ、コンポーネント、時間及び話者のインデックスである。また、 M 、 T 及び S は、それぞれ、コンポーネント、フレーム及び話者の総数である。

30

【 数 9 】

$$\hat{\mu}_m^{(s)} \text{ 及び } \hat{\Sigma}_m^{(s)}$$

【 0 0 8 0 】

の正確な形は、適用される話者依存変換のタイプに依存する。

【 0 0 8 1 】

最も一般的な方法において、話者依存変換は、以下を含む。

【 数 1 0 】

40

— 話者依存加重値のセット $\lambda_{q(m)}^{(s)}$

— 話者依存クラスター $\mu_{c(m,x)}^{(s)}$

— 線形変換のセット $[A_{r(m)}^{(s)}, b_{r(m)}^{(s)}]$

50

【 0 0 8 2 】

【 数 1 1 】

ステップ 2 1 1 において可能なすべての話者依存変換を適用した後に、話者

s のための確率分布 m の平均ベクトル $\hat{\mu}_m^{(s)}$ 及び共分散マトリックス $\hat{\Sigma}_m^{(s)}$ は、次のようになる。

$$\hat{\mu}_m^{(s)} = \mathbf{A}_{r(m)}^{(s)-1} \left(\sum_i \lambda_i^{(s)} \mu_{c(m,i)} + \left(\mu_{c(m,x)}^{(s)} - \mathbf{b}_{r(m)}^{(s)} \right) \right) \quad \dots \text{式 (6)} \quad 10$$

$$\hat{\Sigma}_m^{(s)} = \left(\mathbf{A}_{r(m)}^{(s)\top} \Sigma_{v(m)}^{-1} \mathbf{A}_{r(m)}^{(s)} \right)^{-1} \quad \dots \text{式 (7)} \quad 20$$

【 0 0 8 3 】

ここで、 $\mu_{c(m,i)}$ は、式 (1) に記述されるように、コンポーネント m のためのクラスター I の平均であり、 $\mu_{c(m,x)}^{(s)}$ は、以下で説明する、話者 s のための追加のクラスターのコンポーネント m のための平均ベクトルである。

【 0 0 8 4 】

$\mathbf{A}_{r(m)}^{(s)}$ 及び $\mathbf{B}_{r(m)}^{(s)}$ は、話者 s のための回帰クラス r (m) に関連する線形変換行列及びバイアス・ベクトルである。

【 0 0 8 5 】

R は、回帰クラスの総数であり、 $r(m) \in \{1, \dots, R\}$ は、コンポーネント m が属する回帰クラスを示す。 30

【 0 0 8 6 】

いかなる一次変換も適用されないならば、 $\mathbf{A}_{r(m)}^{(s)}$ 及び $\mathbf{B}_{r(m)}^{(s)}$ は、それぞれ、単位行列及びゼロベクトルになる。

【 0 0 8 7 】

後で説明される理由のために、この実施形態では、複数の共分散は、クラスター化され、複数の決定木に配置される。ここで、 $v(m) \in \{1, \dots, V\}$ は、コンポーネント m の共同分散行列が属する共分散決定木中のリーフ・ノードを表し、V は、分散決定木のリーフ・ノードの総数である。

【 0 0 8 8 】

上記を使用すると、補助関数は、次のように表現することができる。 40

【 数 1 2 】

$$\mathcal{Q}(\mathbf{M}, \mathbf{M}') = -\frac{1}{2} \sum_{m,t,s} \gamma_m(t) \left\{ \log |\hat{\Sigma}_{v(m)}| + (\mathbf{o}(t) - \hat{\mu}_m^{(s)})^T \hat{\Sigma}_{v(m)}^{-1} (\mathbf{o}(t) - \hat{\mu}_m^{(s)}) \right\} + C \quad \dots \text{式 (8)}$$

【 0 0 8 9 】

ここで、C は、M とは独立した定数である。

【 0 0 9 0 】

50

したがって、上記を使用し、式(6)、式(7)及び(8)を上記に代入すると、補助関数は、モデル・パラメータが4つの互いに異なる部分に分割され得ることを示す。

【0091】

最初の部分は、規範的モデルのパラメータ(つまり、話者非依存平均 $\{\mu_n\}$ 及び話者非依存共分散 $\{\Sigma_k\}$)である。上記のインデックス n 及び k は、後で説明される平均及び分散決定木のリーフ・ノードを示す。

【0092】

第2の部分は、次の話者依存加重値である。

【数13】

$$\{\lambda_i^{(s)}\}_{s,i}$$

10

【0093】

ここで、 s は話者を示し、 i は、クラスター・インデックス・パラメータを示す。

【0094】

第3の部分は、話者依存クラスター $\mu_{c(m,x)}$ の平均であり、第4の部分は、次のCMLLR制約付き最尤線形回帰変換である。

【数14】

$$\{\mathbf{A}_d^{(s)}, \mathbf{b}_d^{(s)}\}_{s,d}$$

20

【0095】

ここで、 s は話者を示し、 d はコンポーネント又はコンポーネント m が属する話者回帰クラスを示す。

【0096】

補助関数が上記の方法で表現されれば、それは、話者及び音声特性パラメータのML値、話者依存パラメータのML値、並びに、音声特性依存パラメータのML値を得るために、各々の変数に関して順に最大化される。

【0097】

詳しくは、平均のML推定を決定するために、下記手続きが実行される。

【0098】

以下の方程式を単純化するために、線形変換が適用されないものと仮定する。線形変換が適用されるならば、オリジナルの観測ベクトル $\{\mathbf{o}_r(t)\}$ は、次の変換ベクトルにより置き換えられる必要がある。

【数15】

$$\{\hat{\mathbf{o}}_{r(m)}^{(s)}(t) = \mathbf{A}_{r(m)}^{(s)} \mathbf{o}(t) + \mathbf{b}_{r(m)}^{(s)}\} \quad \dots \text{式(9)}$$

40

【0099】

同様に、追加のクラスターは存在しないものと仮定する。

【0100】

訓練の間に追加のクラスターを含むことは、 $\mathbf{A}_{r(m)}^{(s)}$ が単位行列であり且つ

【数16】

$$\{\mathbf{b}_{r(m)}^{(s)} = \boldsymbol{\mu}_{c(m,x)}^{(s)}\}$$

【0101】

50

である線形変換を付加することにちょうど等しい。

【 0 1 0 2 】

最初に、式 (4) の補助関数が、以下のように μ_n で微分される。

【 数 1 7 】

$$\frac{\partial Q(\mathcal{M}; \hat{\mathcal{M}})}{\partial \mu_n} = k_n - G_{nn} \mu_n - \sum_{\nu \neq n} G_{n\nu} \mu_\nu$$

10

…式 (1 0)

【 0 1 0 3 】

ここで、

【 数 1 8 】

$$G_{n\nu} = \sum_{\substack{m, i, j \\ c(m, i) = n \\ c(m, j) = \nu}} G_{ij}^{(m)}, \quad k_n = \sum_{\substack{m, i \\ c(m, i) = n}} k_i^{(m)}.$$

20

…式 (1 1)

【 0 1 0 4 】

である。

【 0 1 0 5 】

$G^{(m)}_{ij}$ 及び $k^{(m)}_i$ は、蓄積された統計データ (accumulated statistics) である。

【 数 1 9 】

$$G_{ij}^{(m)} = \sum_{t, s} \gamma_m(t, s) \lambda_{i, q(m)}^{(s)} \Sigma_{v(m)}^{-1} \lambda_{j, q(m)}^{(s)}$$

$$k_i^{(m)} = \sum_{t, s} \gamma_m(t, s) \lambda_{i, q(m)}^{(s)} \Sigma_{v(m)}^{-1} o(t).$$

30

…式 (1 2)

【 0 1 0 6 】

導関数を 0 にセットして法線方向において式を最大化することによって、 μ_n の ML 推定、すなわち、

40

【 数 2 0 】

$$\hat{\mu}_n$$

【 0 1 0 7 】

について次の式が得られる。

【数 2 1】

$$\hat{\mu}_n = G_{nn}^{-1} \left(k_n - \sum_{\nu \neq n} G_{n\nu} \mu_\nu \right)$$

…式 (1 3)

10

【0 1 0 8】

μ_n の M L 推定はまた、 μ_k に依存することに留意されるべきである（ここで、 k は n と等しくない）。インデックス n は、平均ベクトルの判定木のリーフ・ノードを表わすために用いられるのに対して、インデックス k は、共分散決定木のリーフ・ノードを表わす。したがって、収束するまですべての μ_n にわたり繰り返すことによって最適化を実行することが必要である。

【0 1 0 9】

これは、次式を解くことによりすべての μ_n を同時に最適化することによって実行することができる。

20

【数 2 2】

$$\begin{bmatrix} G_{11} & \dots & G_{1N} \\ \vdots & \ddots & \vdots \\ G_{N1} & \dots & G_{NN} \end{bmatrix} \begin{bmatrix} \hat{\mu}_1 \\ \vdots \\ \hat{\mu}_N \end{bmatrix} = \begin{bmatrix} k_1 \\ \vdots \\ k_N \end{bmatrix},$$

…式 (1 4)

30

【0 1 1 0】

しかしながら、訓練データが小さいか又は N が非常に大きい場合、式 (7) の係数行列はフルランクを有することができない。この問題は、特異値分解又は他の良く知られた行列因数分解技術を用いることにより回避することができる。

【0 1 1 1】

その後、同じプロセスが、共分散の M L 推定を実行するために行われる。つまり、式 (8) に示される補助関数が Σ_k で微分され、次式が与えられる。

【数 2 3】

$$\hat{\Sigma}_k = \frac{\sum_{v(m)=k} t,s,m \gamma_m(t,s) \bar{o}(t) \bar{o}(t)^\top}{\sum_{v(m)=k} t,s,m \gamma_m(t,s)}$$

40

…式 (1 5)

【0 1 1 2】

ここで、

【数 2 4】

$$\bar{o}(t) = o(t) - \mu_m^{(s)} \quad \dots \text{式 (16)}$$

【0 1 1 3】

である。

【0 1 1 4】

話者依存加重値及び話者依存線形変換のためのML推定も、同じ方法で、つまり、ML推定が求められるパラメータに関して補助関数を微分し、そして、微分の値を0にセットすることで、得ることができる。

10

【0 1 1 5】

話者依存加重値のために、これは次を与える。

【数 2 5】

$$\lambda_q^{(s)} = \left(\sum_{\substack{t,m \\ q(m)=q}} \gamma_m(t,s) M_m^\top \Sigma^{-1} M_m \right)^{-1} \sum_{\substack{t,m \\ q(m)=q}} \gamma_m(t,s) M_m^\top \Sigma^{-1} o(t) \quad \dots \text{式 (17)}$$

20

【0 1 1 6】

一つの実施形態において、該プロセスは反復する方法で、実行される。図11のフローチャートを参照して、この基本的なシステムが説明される。

30

【0 1 1 7】

ステップS301において、複数のオーディオ音声の入力が受信される。この実例となる例において、4人の話者が使用される。

【0 1 1 8】

次に、ステップS303において、4人の音声のそれぞれごとに、音響モデルが訓練され生成される。この実施形態において、4つのモデルのそれぞれは、一つの音声からのデータを使用して訓練されるだけである。

【0 1 1 9】

クラスター適応可能なモデルは、次のように、初期化され訓練される。ステップS305において、クラスターPの個数が、V+1にセットされる。ここで、Vは、音声の個数(4)である。

40

【0 1 2 0】

ステップS307において、一つのクラスター(クラスター1)が、バイアス・クラスターとして決定される。バイアス・クラスターのための各決定木と、関連する各クラスター平均ベクトルは、ステップS303において最良のモデルを生成した音声を使用して初期化される。この例では、各々の音声は、タグ「音声A」、「音声B」、「音声C」及び「音声D」を与えられる。ここで、音声Aが最良のモデルを生成したものと仮定する。共分散マトリックス、マルチ空間確率分布(multi-space probability distributions)(

50

M S D)に関する空間加重値、及び、構造を共有しているそれらのパラメータもまた、音声 A モデルのそれらに初期化される。

【 0 1 2 1 】

各々の二分決定木は、すべてのコンテキストを表現する単一のルート・ノードから始まる局所的最適化法で構築される。この実施形態において、コンテキストによって、次のベース（音声ベース、言語ベース、及び、韻律ベース）が使用される。各々のノードが作成されるとともに、コンテキストに関する次の最適な質問が選択される。いずれの質問が尤度の最大の増加をもたらすか及び訓練例において生成される終端ノードに基づいて、質問が選択される。

【 0 1 2 2 】

その後、訓練データに総尤度の最大の増加を提供するために、その最適の質問を用いて分割することができる終端ノードを発見するために、終端ノードのセットが検索される。この増加が閾値を越えるとすれば、該ノードは最適な質問を用いて分割され、2つの新たな終端ノードが作成される。更に分割しても、尤度分割に適用される閾値を越えないことにより、新たな終端ノードを形成することができない場合、そのプロセスは停止する。

【 0 1 2 3 】

このプロセスは例えば図 1 2 に示される。平均決定木中の第 n 番目の終端ノードは、質問 q により 2 の新たな終端ノード n_+^q 及び n_-^q に分割される。この分割により達成される尤度の増加は、以下のように計算することができる。

【 数 2 6 】

$$\begin{aligned} \mathcal{L}(n) = & -\frac{1}{2}\mu_n^\top \left(\sum_{m \in \mathcal{S}(n)} G_{ii}^{(m)} \right) \mu_n \\ & + \mu_n^\top \sum_{m \in \mathcal{S}(n)} \left(k_i^{(m)} - \sum_{j \neq i} G_{ij}^{(m)} \mu_{c(m,j)} \right) \end{aligned} \quad \dots \text{式 (18)}$$

【 0 1 2 4 】

$\mathcal{S}(n)$ は、ノード n に関連するコンポーネントのセットを示す。 μ_n に関して不変である項は含まれない点に留意されるべきである。

【 0 1 2 5 】

ここで、 c は、 μ_n とは独立した定数項である。 μ_n の最大尤度は式 (1 3) により与えられる。それゆえ、上記は、次のように書くことができる。

【数 2 7】

$$\mathcal{L}(n) = \frac{1}{2} \hat{\mu}_n^\top \left(\sum_{m \in \mathcal{S}(n)} G_{ii}^{(m)} \right) \hat{\mu}_n$$

…式 (19)

10

【0126】

したがって、ノード n を n_+^q 及び n_-^q へ分割することにより得られる尤度は、次式により与えられる。

【数 2 8】

$$\Delta \mathcal{L}(n; q) = \mathcal{L}(n_+^q) + \mathcal{L}(n_-^q) - \mathcal{L}(n)$$

…式 (20)

20

【0127】

したがって、上記を使用して、各々のクラスターの決定木を構築することは可能である。ここで、木は、最初に木において最適な質問が尋ねられ、分割の尤度に従う階層の順に決定が配列されるように、配列される。その後、加重値が各々のクラスターに適用される。

【0128】

決定木は、同様に、分散のために構築され得る。共分散決定木は、以下のように構築される：共分散決定木中のケース終端ノードが、質問 q により 2 の新たな終端ノード k_+^q 及び k_-^q に分割されるならば、クラスター分散行列及び分割による増加は、以下のように表現される。

30

【数 2 9】

$$\Sigma_k = \frac{\sum_{\substack{m, t, s \\ v(m)=k}} \gamma_m(t) \Sigma_{v(m)}}{\sum_{\substack{m, t, s \\ v(m)=k}} \gamma_m(t)}$$

40

…式 (21)

【0129】

【数 3 0】

$$\mathcal{L}(k) = -\frac{1}{2} \sum_{\substack{m, t, s \\ v(m)=k}} \gamma_m(t) \log |\Sigma_k| + D$$

…式 (22)

10

【0130】

ここで、Dは、 $\{\Sigma_k\}$ とは独立した定数である。したがって、尤度の増加は、次のようになる。

【数 3 1】

$$\Delta \mathcal{L}(k, q) = \mathcal{L}(k_+^q) + \mathcal{L}(k_-^q) - \mathcal{L}(k)$$

…式 (23)

20

【0131】

ステップS309において、クラスター2, ..., Pの各々に特定の音声タグ (voice tag) が割り当てられる。例えば、クラスター2, 3, 4及び5が、それぞれ話者B、C、D及びAに対する。音声Aがバイアス・クラスターを初期化するのに用いられたので、それは初期化されるべき最後のクラスターに割り当てられることに留意されるべきである。ステップS311において、CAT補間加重値のセットは、割り当てられた音声タグに従って、以下のように単に1又は0にセットされる。

【数 3 2】

$$\lambda_i^{(s)} = \begin{cases} 1.0 & i=0 \text{ のとき} \\ 1.0 & \text{音声タグ}(s)=i \text{ のとき} \\ 0.0 & \text{それ以外} \end{cases}$$

30

【0132】

この具体例では、ストリームごと話者ごとに大域的な加重値 (global weights) が存在する。話者/ストリームの組み合わせごとに、3セットの加重値がセットされる：無音、音声及びポーズについて。

【0133】

40

ステップS313において、各々のクラスター2, ..., (P-1)について順番に、以下のようにクラスターが初期化される。関連する音声 (voice) のための音声データ (例えば、クラスター2のための音声B) は、ステップS303で訓練される関連する音声のための単一話者モデルを使用して、調整 (aligned) される。これらの調整を所与として、統計値が計算され、そして、クラスターのための決定木及び平均値が推定される。クラスターのための平均値は、ステップS311でセットされた加重値を使用して、クラスター平均の正規化された加重和として、計算される。すなわち、実際には、これは、所与のコンテキストに関する平均値 (そのコンテキストに関するバイアス・クラスターの平均の加重和 (いずれの場合も加重値1) である)、そして、クラスター2におけるそのコンテキストに関する音声Bモデルの平均をもたらす。

50

【 0 1 3 4 】

それから、ステップ S 3 1 5 において、全 4 つの音声からのすべてのデータを使用して、バイアス・クラスターのために決定木が再構築され、関連する平均及び分散パラメータが再推定される。

【 0 1 3 5 】

音声 B、C 及び D のためのクラスターを加えた後に、バイアス・クラスターは、同時に全 4 つの音声を使用して、再推定される。

【 0 1 3 6 】

ステップ S 3 1 7 において、クラスター P (音声 A) は、ステップ S 3 1 3 で説明されるように、音声 A だけからのデータを使用して、他のクラスターに関して、初期化される。

10

【 0 1 3 7 】

クラスターが上記のように初期化されたならば、その後、C A T モデルは、以下のように、更新され / 訓練される。

【 0 1 3 8 】

ステップ S 3 1 9 において、C A T 加重値が固定された状態で、クラスター 1 から P まで、1 クラスターずつ、決定木が再構成される。ステップ S 3 2 1 において、新たな平均及び分散が C A T モデルで推定される。次に、ステップ S 3 2 3 において、各クラスターごとに、新たな C A T 加重値が推定される。一つの実施形態では、該プロセスは、収束するまで S 3 2 1 ヘループバックする。パラメータ及び加重値は、上記パラメータのより良い推定を得るために、パウム・ウェルチ・アルゴリズムの補助関数を用いて実行される最尤計算を使用して、推定される。

20

【 0 1 3 9 】

前に説明されたように、パラメータは反復プロセスにより推定される。

【 0 1 4 0 】

更なる実施形態では、ステップ S 3 2 3 において、プロセスは、収束するまで各々の繰り返しの間に決定木が再構成されるように、ステップ S 3 1 9 ヘループバックする。

【 0 1 4 1 】

更なる実施形態では、前述のような話者依存変換が使用される。ここでは、該変換が適用されるように、ステップ S 3 2 3 の後、話者依存変換が挿入され、それから、変換モデルは、収束するまで繰り返される。一つの実施形態では、該変換は、各々の繰り返しにおいて、更新されるであろう。

30

【 0 1 4 2 】

図 1 2 は、決定木の形をとるクラスター 1 ~ P を示す。この単純化された例では、ちょうどクラスター 1 に 4 つの終端ノードが存在し、クラスター P に 3 つの終端ノードが存在する。決定木は対称である必要がない、つまり、各々の決定木が異なる数の終端ノードを有することができることに留意することは重要である。木における終端ノードの数及びブランチの数は、純粹に対数尤度分割によって決定される。対数尤度分割は、最初の決定において最大の分割を達成し、次いで、より大きな分割をもたらす質問の順に質問が尋ねられる。達成された分割が閾値未満ならば、終端ノードの分割は終了する。

40

【 0 1 4 3 】

上記は、実行されるべき次の合成を可能にする規範的モデルを生成する：

1 . 4 つの音声のうちの任意のものが、その音声に対応する加重値ベクトルの最終的なセットを使用して合成されることができる。

2 . ランダムな音声は、加重値ベクトルを任意の位置にセットすることによって、C A T モデルが及び音響空間から合成することができる。

【 0 1 4 4 】

更なる例において、音声特性を合成するために、アシスタントが使用される。ここで、該システムは、同一の特徴をもつ目標音声 (target voice) の入力を与えられる。

【 0 1 4 5 】

50

図 1 3 は、一つの例を示す。最初に、入力目標音声が入力ステップ 5 0 1 で受信される。次に、規範的モデルの加重値（すなわち、前もって訓練されたクラスターの加重値）は、ステップ 5 0 3 で、目標音声にマッチするように調整される。

【 0 1 4 6 】

それから、ステップ S 5 0 3 で得られる新たな加重値を使用して、音声（audio）が出力される。

【 0 1 4 7 】

更なる実施形態では、新たな音声のために新たなクラスターが提供される、より複雑な方法が使用される。これは、図 1 4 を参照して説明される。

【 0 1 4 8 】

図 1 3 のように、最初に、ステップ S 5 0 1 において、目標音声を受信される。加重値は、それから、ステップ S 5 0 3 において、目標音声に最もマッチするように調整される。

【 0 1 4 9 】

それから、ステップ S 5 0 7 において、新たなクラスターが、目標音声のモデルに追加される。次に、図 1 1 を参照して説明された方法と同様な方法で、新たな話者依存クラスターについて、決定木が構築される。

【 0 1 5 0 】

それから、ステップ S 5 1 1 において、新たなクラスターについて、音響モデル・パラメータ（すなわち、この例では、平均）が計算される。

【 0 1 5 1 】

次に、ステップ S 5 1 3 において、すべてのクラスターについて、加重値が更新される。それから、ステップ S 5 1 5 において、新たなクラスターの構造が更新される。

【 0 1 5 2 】

前述のように、ステップ S 5 0 5 において、新たなクラスターをもつ新たな加重値を使用して、新たな目標音声をもつオーディオが出力される。

【 0 1 5 3 】

この実施形態では、これは訓練データが合成時間に利用できることを要求するであろうから、ステップ S 5 1 5 において、他のクラスターはこのときに更新されないことに留意されるべきである。

【 0 1 5 4 】

更なる実施形態では、ステップ S 5 1 5 の後で、各クラスターが更新される。それゆえ、フローチャートは、収束するまでステップ S 5 0 9 ヘループバックする。

【 0 1 5 5 】

最後に、目標話者との類似を更に改善するために、該モデルの上に、例えば C M L L R のような線形変換を適用することができる。この変換の回帰クラスは、大域的であることができ、あるいは、話者依存であることができる。

【 0 1 5 6 】

もう一つのケースでは、回帰クラスの共有構造（tying structure）は、話者依存クラスターの決定木から、又は、話者依存加重値を規範的モデルに適用し、追加のクラスターを加えた後に得られる分布のクラスタリングから、得ることが出来る。

【 0 1 5 7 】

初めは、バイアス・クラスターは、話者 / 音声非依存特性を表し、一方、他のクラスターは、それらの関連する音声データセットを表す。訓練が進むにつれて、音声に対するクラスターの正確な割り当ては、より正確さの低いものになる。クラスター及び C A T 加重値は、幅広い音響空間（broad acoustic space）を表す。

【 0 1 5 8 】

特定の実施形態が説明されたが、これらの実施形態はただ例として示されたものであり、本発明の範囲を制限することが意図されるものではない。実際に、本明細書で説明された新規な方法及び装置は、種々の他の形で実施されても良い；更に、本明細書で説明され

10

20

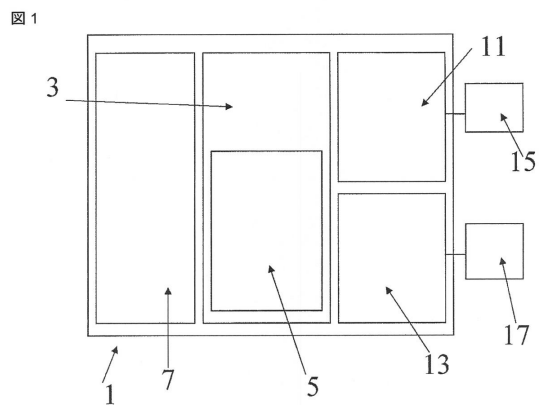
30

40

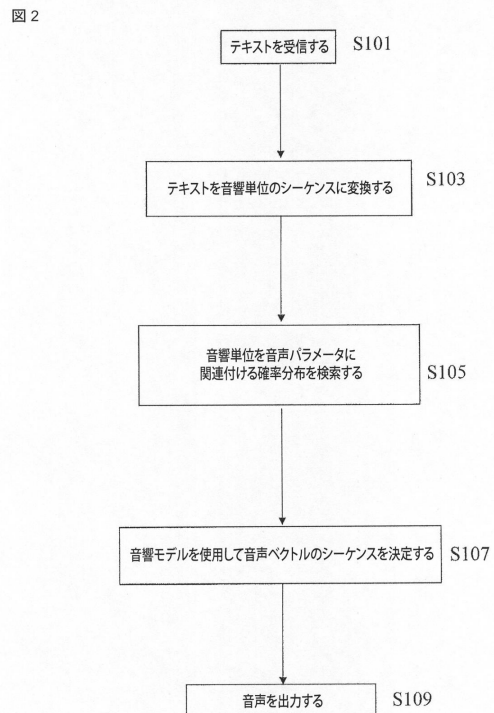
50

た方法及び装置の形における様々な省略、置き換え及び変更は、本発明の精神を逸脱せずになされ得る。添付の特許請求の範囲及びそれらの均等物は、本発明の範囲及び精神に含まれるであろうそのような修正の形をカバーすることが意図される。

【図 1】

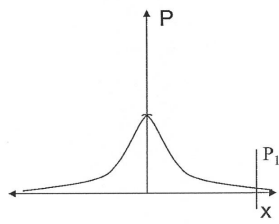


【図 2】



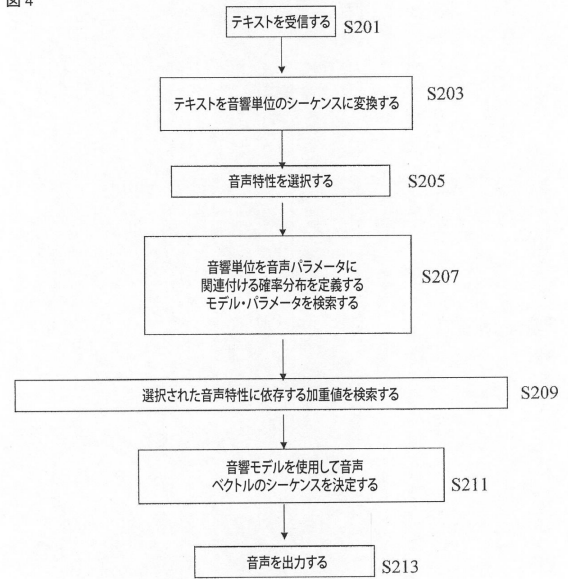
【図 3】

図 3



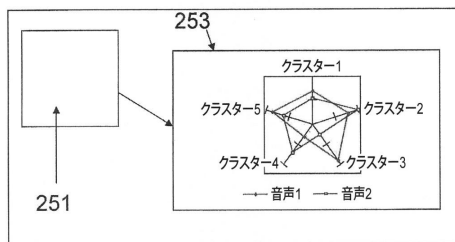
【図 4】

図 4



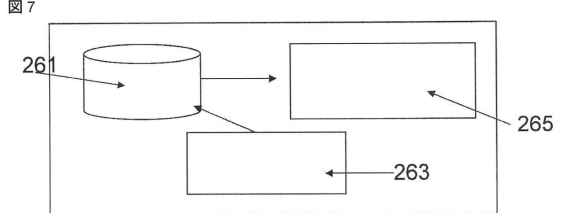
【図 5】

図 5



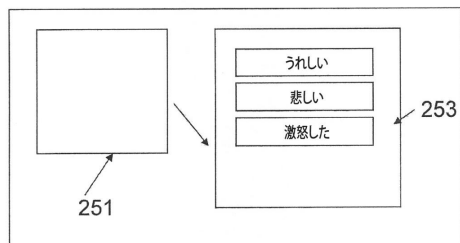
【図 7】

図 7



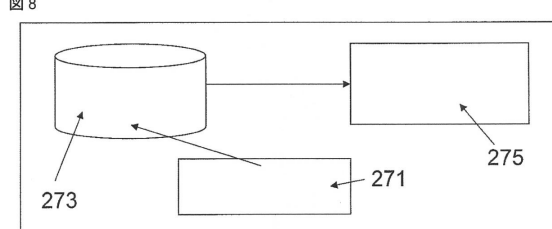
【図 6】

図 6



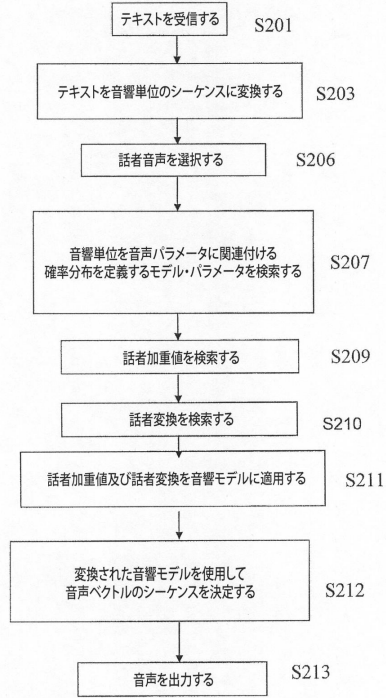
【図 8】

図 8



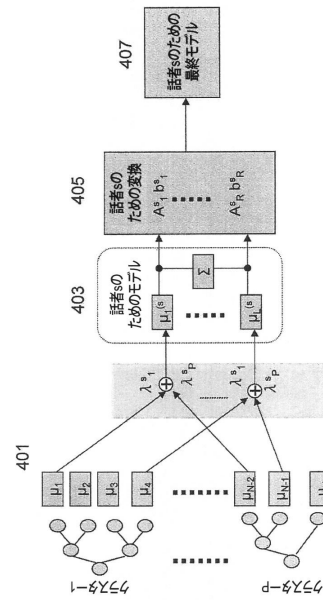
【図 9 a】

図 9a



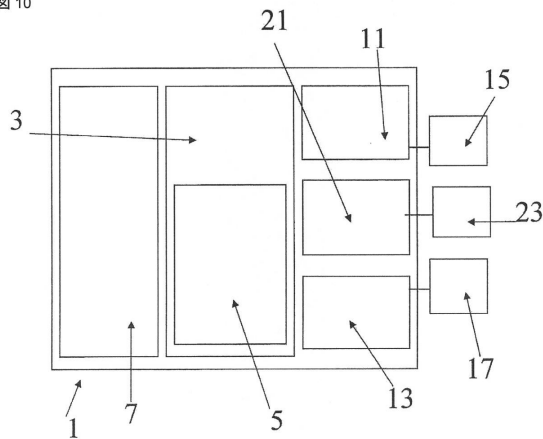
【図 9 b】

図 9b



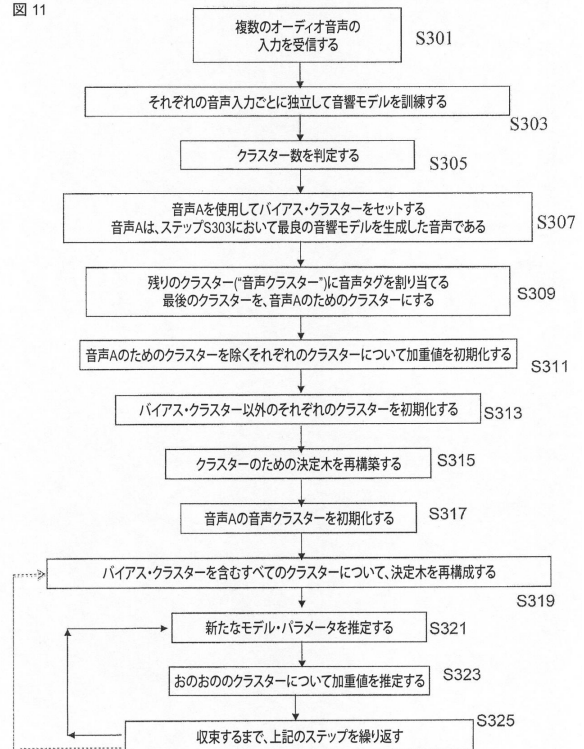
【図 10】

図 10

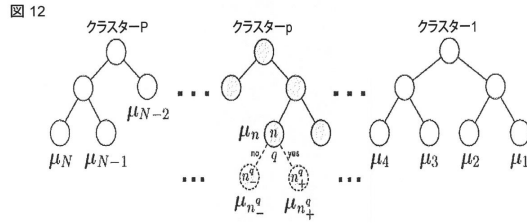


【図 11】

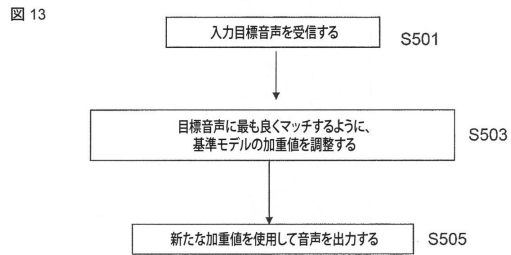
図 11



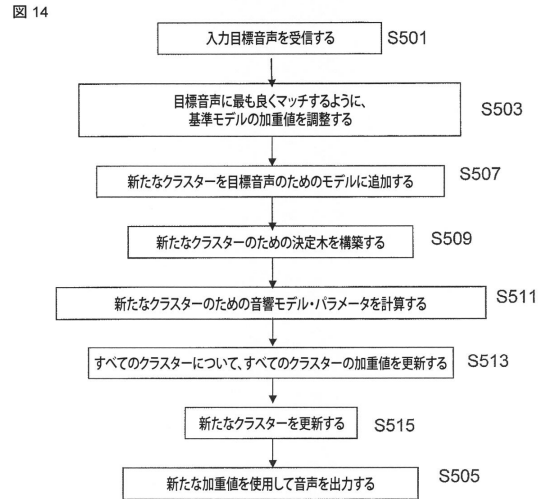
【図 12】



【図 13】



【図 14】



フロントページの続き

(51)Int.Cl.

F I

G 1 0 L 13/10 1 1 4

(74)代理人 100124394

弁理士 佐藤 立志

(74)代理人 100112807

弁理士 岡田 貴志

(74)代理人 100111073

弁理士 堀内 美保子

(72)発明者 赤嶺 政巳

イギリス国、 シービー４・０ジーゼット、 ケンブリッジシャー、 ケンブリッジ、 ミルトン・ロード、 ケンブリッジ・サイエンス・パーク ２０８、 トーシバ・リサーチ・ヨーロッパ・リミテッド、 ケンブリッジ・リサーチ・ラボラトリー内

(72)発明者 ラトレ・マルティネス・ハビエル

イギリス国、 シービー４・０ジーゼット、 ケンブリッジシャー、 ケンブリッジ、 ミルトン・ロード、 ケンブリッジ・サイエンス・パーク ２０８、 トーシバ・リサーチ・ヨーロッパ・リミテッド、 ケンブリッジ・リサーチ・ラボラトリー内

(72)発明者 ワン・ビンセント・ピン・ルン

イギリス国、 シービー４・０ジーゼット、 ケンブリッジシャー、 ケンブリッジ、 ミルトン・ロード、 ケンブリッジ・サイエンス・パーク ２０８、 トーシバ・リサーチ・ヨーロッパ・リミテッド、 ケンブリッジ・リサーチ・ラボラトリー内

(72)発明者 チン・カン・クホン

イギリス国、 シービー４・０ジーゼット、 ケンブリッジシャー、 ケンブリッジ、 ミルトン・ロード、 ケンブリッジ・サイエンス・パーク ２０８、 トーシバ・リサーチ・ヨーロッパ・リミテッド、 ケンブリッジ・リサーチ・ラボラトリー内

(72)発明者 ゲールズ・マーク・ジョン・フランシス

イギリス国、 シービー４・０ジーゼット、 ケンブリッジシャー、 ケンブリッジ、 ミルトン・ロード、 ケンブリッジ・サイエンス・パーク ２０８、 トーシバ・リサーチ・ヨーロッパ・リミテッド、 ケンブリッジ・リサーチ・ラボラトリー内

(72)発明者 ニル・キャサリン・マリー

イギリス国、 シービー４・０ジーゼット、 ケンブリッジシャー、 ケンブリッジ、 ミルトン・ロード、 ケンブリッジ・サイエンス・パーク ２０８、 トーシバ・リサーチ・ヨーロッパ・リミテッド、 ケンブリッジ・リサーチ・ラボラトリー内

(72)発明者 チュン・ビュン・ハ

イギリス国、 シービー４・０ジーゼット、 ケンブリッジシャー、 ケンブリッジ、 ミルトン・ロード、 ケンブリッジ・サイエンス・パーク ２０８、 トーシバ・リサーチ・ヨーロッパ・リミテッド、 ケンブリッジ・リサーチ・ラボラトリー内

審査官 上田 雄

(56)参考文献 特開２００７－１８３４２１（ＪＰ，Ａ）

特開２００３－２３３３８８（ＪＰ，Ａ）

特開２００３－３３７５９２（ＪＰ，Ａ）

特開平０９－１３８７６７（ＪＰ，Ａ）

特開２００３－１７７７７２（ＪＰ，Ａ）

国際公開第２０１０／１４２９２８（ＷＯ，Ａ１）

特開２０１１－０２８１３１（ＪＰ，Ａ）

(58)調査した分野(Int.Cl. , D B 名)

G 1 0 L 1 3 / 0 0 - 1 3 / 1 0