

(19) 日本国特許庁(JP)

(12) 公開特許公報(A)

(11) 特許出願公開番号

特開2010-54733

(P2010-54733A)

(43) 公開日 平成22年3月11日(2010.3.11)

(51) Int. Cl.		F I			テーマコード (参考)	
G 1 0 L	15/04	(2006.01)	G 1 0 L	15/04	3 0 0 D	5 D 0 1 5
G 1 0 L	15/28	(2006.01)	G 1 0 L	15/04	3 0 0 B	
G 1 0 L	17/00	(2006.01)	G 1 0 L	15/28	4 0 0	
			G 1 0 L	17/00	2 0 0 B	

審査請求 未請求 請求項の数 12 O L (全 18 頁)

(21) 出願番号	特願2008-218677 (P2008-218677)	(71) 出願人	000004226
(22) 出願日	平成20年8月27日 (2008. 8. 27)		日本電信電話株式会社
			東京都千代田区大手町二丁目3番1号
		(74) 代理人	100121706
			弁理士 中尾 直樹
		(74) 代理人	100066153
			弁理士 草野 卓
		(74) 代理人	100128705
			弁理士 中村 幸雄
		(72) 発明者	荒木 章子
			東京都千代田区大手町二丁目3番1号 日
			本電信電話株式会社内
		(72) 発明者	石塚 健太郎
			東京都千代田区大手町二丁目3番1号 日
			本電信電話株式会社内

最終頁に続く

(54) 【発明の名称】 複数信号区間推定装置、複数信号区間推定方法、そのプログラムおよび記録媒体

(57) 【要約】

【課題】 音声の収録中に話者位置の移動が生じてても、同一話者には同一インデックスを付与することを可能とする。

【解決手段】 周波数領域変換部 110 が観測信号を所定長のフレームに順次切り出して当該フレームごとに周波数領域に変換し、音声区間推定部 120 が周波数領域の観測信号に基づき、各フレームが音声区間に該当するかどうかを推定し、到来方向推定部 130 が周波数領域の観測信号に基づき、当該周波数領域の観測信号の到来方向を各フレームごとに推定し、到来方向分類部 140 が音声区間に該当すると推定された各フレームを、到来方向の類似性に基づき話者ごとのクラスタに分類する。そして、話者同定部 250 が所定の時刻までに同一クラスタに分類された各フレームの周波数領域の観測信号に基づき、当該クラスタに係る話者のモデルをクラスタごとに作成し、当該所定の時刻以降の観測信号の話者を各話者のモデルに基づき推定する。

【選択図】 図 1

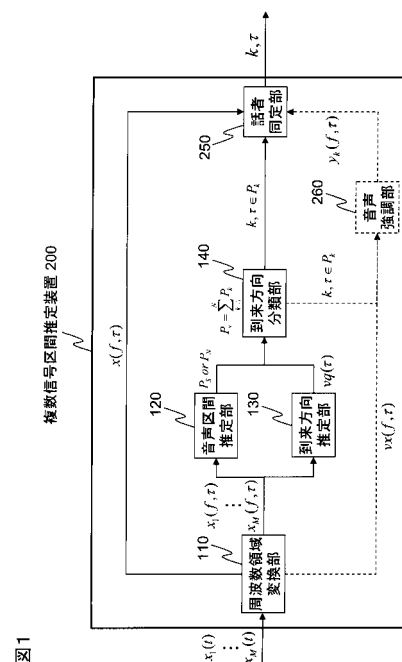


図 1

【特許請求の範囲】**【請求項 1】**

複数のマイクによりそれぞれ収録された、複数の話者による発話音声が含まれる観測信号から、それぞれの話者の発話区間を推定する複数信号区間推定装置であって、

上記観測信号を所定長のフレームに順次切り出し、当該フレームごとに周波数領域に変換する周波数領域変換部と、

周波数領域に変換された上記観測信号（以下、「周波数領域観測信号」という）に基づき、各フレームが音声区間に該当するか否かを推定する音声区間推定部と、

上記周波数領域観測信号に基づき、当該周波数領域観測信号の到来方向を各フレームごとに推定する到来方向推定部と、

上記音声区間に該当すると推定された各フレームを、上記到来方向の類似性に基づき上記話者ごとのクラスタに分類する到来方向分類部と、

所定の時刻までに同一クラスタに分類された各フレームの上記周波数領域観測信号に基づき、当該クラスタに係る上記話者のモデルをクラスタごとに作成し、当該所定の時刻以降の上記観測信号の話者を、各話者のモデルに基づき推定する話者同定部と、
を備えることを特徴とする複数信号区間推定装置。

10

【請求項 2】

請求項 1 に記載の複数信号区間推定装置において、

上記話者同定部は、

上記周波数領域観測信号の各フレームの音声特徴量を計算する特徴抽出手段と、

20

上記所定の時刻までに同一クラスタに分類された各フレームの音声特徴量を用いて、当該クラスタに係る上記話者のモデルを各クラスタごとに作成して出力するとともに、上記所定の時刻までの各フレームのインデックスと当該各フレームが属するクラスタに係る話者のインデックスとの組を出力するモデル学習手段と、

上記所定の時刻以降に同一クラスタに分類された互いに接続されたフレーム（以下、「セグメント」という）の音声特徴量について、上記話者のモデルに対する尤度を各モデルごとに計算し、最大尤度をとるモデルに係る話者のインデックスを当該セグメントに付与して、当該セグメントに含まれる各フレームのインデックスとともに出力する尤度計算手段と、

を備えることを特徴とする複数信号区間推定装置。

30

【請求項 3】

複数のマイクによりそれぞれ収録された、複数の話者による発話音声が含まれる観測信号から、それぞれの話者の発話区間を推定する複数信号区間推定装置であって、

上記観測信号を所定長のフレームに順次切り出し、当該フレームごとに周波数領域に変換する周波数領域変換部と、

周波数領域に変換された上記観測信号（以下、「周波数領域観測信号」という）に基づき、各フレームが音声区間に該当するか否かを推定する音声区間推定部と、

上記周波数領域観測信号に基づき、当該周波数領域観測信号の到来方向を各フレームごとに推定する到来方向推定部と、

上記音声区間に該当すると推定された各フレームを、上記到来方向の類似性に基づき上記話者ごとのクラスタに分類する到来方向分類部と、

40

上記周波数領域観測信号に基づき、上記クラスタに係る上記話者ごとに強調した信号（以下、「強調信号」という）を生成する音声強調部と、

所定の時刻までの上記強調信号に基づき、上記話者のモデルを話者ごとに作成し、当該所定の時刻以降の上記観測信号の話者を、各話者のモデルに基づき推定する話者同定部と、
を備えることを特徴とする複数信号区間推定装置。

【請求項 4】

請求項 3 に記載の複数信号区間推定装置において、

上記話者同定部は、

50

上記強調信号の各フレームの音声特徴量を計算する特徴抽出手段と、

上記所定の時刻までの上記強調信号の各フレームの音声特徴量を用いて、上記話者のモデルを各話者ごとに作成して出力するとともに、上記所定の時刻までの各フレームのインデックスと当該各フレームが属するクラスに属する話者のインデックスとの組を出力するモデル学習手段と、

上記所定の時刻以降に同一クラスに分類された互いに接続されたフレーム（以下、「セグメント」という）の音声特徴量について、上記話者のモデルに対する尤度を各モデルごとに計算し、最大尤度をとるモデルに係る話者のインデックスを当該セグメントに付与して、当該セグメントに含まれる各フレームのインデックスとともに出力する尤度計算手段と、

10

を備えることを特徴とする複数信号区間推定装置。

【請求項 5】

複数のマイクによりそれぞれ収録された、複数の話者による発話音声が含まれる観測信号から、それぞれの話者の発話区間を推定する複数信号区間推定方法であって、

上記観測信号を所定長のフレームに順次切り出し、当該フレームごとに周波数領域に変換する周波数領域変換ステップと、

周波数領域に変換された上記観測信号（以下、「周波数領域観測信号」という）に基づき、各フレームが音声区間に該当するか否かを推定する音声区間推定ステップと、

上記周波数領域観測信号に基づき、当該周波数領域観測信号の到来方向を各フレームごとに推定する到来方向推定ステップと、

20

上記音声区間に該当すると推定された各フレームを、上記到来方向の類似性に基づき上記話者ごとのクラスに分類する到来方向分類ステップと、

所定の時刻までに同一クラスに分類された各フレームの上記周波数領域観測信号に基づき、当該クラスに係る上記話者のモデルをクラスごとに作成し、当該所定の時刻以降の上記観測信号の話者を、各話者のモデルに基づき推定する話者同定ステップと、
を実行することを特徴とする複数信号区間推定方法。

【請求項 6】

請求項 5 に記載の複数信号区間推定装置において、

上記話者同定ステップは、

上記周波数領域観測信号の各フレームの音声特徴量を計算する特徴抽出サブステップと

30

、
上記所定の時刻までに同一クラスに分類された各フレームの音声特徴量を用いて、当該クラスに係る上記話者のモデルを各クラスごとに作成して出力するとともに、上記所定の時刻までの各フレームのインデックスと当該各フレームが属するクラスに係る話者のインデックスとの組を出力するモデル学習サブステップと、

上記所定の時刻以降に同一クラスに分類された互いに接続されたフレーム（以下、「セグメント」という）の音声特徴量について、上記話者のモデルに対する尤度を各モデルごとに計算し、最大尤度をとるモデルに係る話者のインデックスを当該セグメントに付与して、当該セグメントに含まれる各フレームのインデックスとともに出力する尤度計算サブステップと、

40

を実行することを特徴とする複数信号区間推定方法。

【請求項 7】

複数のマイクによりそれぞれ収録された、複数の話者による発話音声が含まれる観測信号から、それぞれの話者の発話区間を推定する複数信号区間推定方法であって、

上記観測信号を所定長のフレームに順次切り出し、当該フレームごとに周波数領域に変換する周波数領域変換ステップと、

周波数領域に変換された上記観測信号（以下、「周波数領域観測信号」という）に基づき、各フレームが音声区間に該当するか否かを推定する音声区間推定ステップと、

上記周波数領域観測信号に基づき、当該周波数領域観測信号の到来方向を各フレームごとに推定する到来方向推定ステップと、

50

上記音声区間に該当すると推定された各フレームを、上記到来方向の類似性に基づき上記話者ごとのクラスタに分類する到来方向分類ステップと、

上記周波数領域観測信号に基づき、上記話者ごとに強調した信号（以下、「強調信号」という）を生成する音声強調ステップと、

所定の時刻までの上記強調信号に基づき、上記話者のモデルを話者ごとに作成し、当該所定の時刻以降の上記観測信号の話者を、各話者のモデルに基づき推定する話者同定ステップと、

を実行することを特徴とする複数信号区間推定方法。

【請求項 8】

請求項 7 に記載の複数信号区間推定方法において、

10

上記話者同定ステップは、

上記強調信号の各フレームの音声特徴量を計算する特徴抽出サブステップと、

上記所定の時刻までの上記強調信号の各フレームの音声特徴量を用いて、上記話者のモデルを話者ごとに作成して出力するとともに、上記所定の時刻までの各フレームのインデックスと当該各フレームが属するクラスタに係る話者のインデックスとの組を出力するモデル学習サブステップと、

上記所定の時刻以降に同一クラスタに分類された互いに接続されたフレーム（以下、「セグメント」という）の音声特徴量について、上記話者のモデルに対する尤度を各モデルごとに計算し、最大尤度をとるモデルに係る話者のインデックスを当該セグメントに付与して、当該セグメントに含まれる各フレームのインデックスとともに出力する尤度計算サブステップと、

20

を実行することを特徴とする複数信号区間推定方法。

【請求項 9】

請求項 6 又は 8 のいずれかに記載の複数信号区間推定方法において、

更に、上記尤度計算サブステップにて上記セグメントに話者のインデックスを付与した後、そのセグメントに属する各フレームの音声特徴量に基づき改めて当該話者のモデルを作成して、当該話者のモデルを更新するモデル更新ステップ

を実行することを特徴とする複数信号区間推定方法。

【請求項 10】

請求項 6、8 又は 9 のいずれかに記載の複数信号区間推定方法において、

30

更に、計算した上記最大尤度が所定の閾値より小さい場合に、新たな話者が参加したと判断し、当該新たな話者のインデックスを上記セグメントに付与するとともに、そのセグメントに属する各フレームの音声特徴量に基づき当該新たな話者のモデルを作成するモデル追加ステップ

を実行することを特徴とする複数信号区間推定方法。

【請求項 11】

請求項 1～4 のいずれかに記載した装置としてコンピュータを機能させるためのプログラム。

【請求項 12】

請求項 11 に記載したプログラムを記録したコンピュータが読み取り可能な記録媒体。

40

【発明の詳細な説明】

【技術分野】

【0001】

本発明は、信号処理の技術分野に属する。特に、複数人の音声信号が混在している音響データについて、各人の音声信号が発せられている区間を推定する複数信号区間推定装置、複数信号区間推定方法、そのプログラムおよび記録媒体に関する。

【背景技術】

【0002】

複数人による会話などを複数のマイクで収録し、「いつ、誰が話したか」を推定する音声区間検出技術は、例えば会議録自動作成において、各発言に発話者を自動的に付与した

50

り、会議収録データに話者情報を付与して録音データの検索や頭出しを容易にしたりする際に有用である。

【 0 0 0 3 】

従来の音声区間検出技術としては、例えば特許文献 1 や非特許文献 1 などが開示されている方法が挙げられる。図 1 1 に従来技術による複数信号区間推定装置 1 0 0 の機能構成例を、図 1 2 にその処理フロー例を示す。複数信号区間推定装置 1 0 0 は、周波数領域変換部 1 1 0 と音声区間推定部 1 2 0 と到来方向推定部 1 3 0 と到来方向分類部 1 4 0 とから構成される。

【 0 0 0 4 】

周波数領域変換部 1 1 0 は、M 本のマイクによりそれぞれ収録した時間領域の観測信号 $x_j(t)$ ($j = 1, \dots, M$) を、例えば 32 ms ごとに窓関数で切り出して（切り出した 1 区間を以下、「フレーム」という）、切り出した各フレーム（インデックスを τ とする）についてフーリエ変換等によりそれぞれ周波数領域の観測信号 $x_j(f, \tau)$ ($f = 1, \dots, L$) に変換する（S 1）。

【 0 0 0 5 】

音声区間推定部 1 2 0 は、周波数領域変換部 1 1 0 で周波数領域に変換された観測信号の各フレームに音声が存在するか否かを、音声存在確率を計算することにより推定する（S 2）。音声存在確率の計算に際しては、例えば非特許文献 2、非特許文献 3 に記載された方法が利用できる。前者で説明すると、該当フレームにおける音声存在確率 $p_v(\tau)$ を次式により求める。

【数 1】

$$p_v(\tau) = \frac{\Lambda(\tau)}{1 + \Lambda(\tau)} \quad (1)$$

$$\log \Lambda(\tau) = \frac{1}{L} \sum_{f=0}^{L-1} \{\gamma(f, \tau) - \log \gamma(f, \tau) - 1\}$$

$$\gamma(f, \tau) = \frac{|x(f, \tau)|^2}{\lambda_N(f)}$$

ここで、 $\lambda_N(f)$ は周波数 f におけるノイズの平均パワー（音声明らかに存在しない録音ファイルの冒頭区間などで求める）、 $x(f, \tau)$ は M 本のマイクにおける周波数領域の観測信号 $x_1(f, \tau) \sim x_M(f, \tau)$ の中から任意に選んだいずれか 1 本についての周波数領域の観測信号である。なお、 $x(f, \tau)$ はすべてのマイクの振幅の平均値として次のように求めても構わない。

【数 2】

$$x(f, \tau) = \frac{1}{M} \sum_{j=1}^M |x_j(f, \tau)| \quad (2)$$

音声区間推定部 1 2 0 は、式(1)により求めた音声存在確率 $p_v(\tau)$ をそのまま出力してもよいし、 $p_v(\tau)$ がある閾値より大きければそのフレームは音声区間 P_S であると判定し、小さければ非音声（ノイズ）区間 P_N と判定して結果を出力してもよい。

【 0 0 0 6 】

到来方向推定部 1 3 0 は、周波数領域変換部 1 1 0 で周波数領域に変換された観測信号の到来方向を各フレームごと又は各フレームの各周波数成分ごとに推定する（S 3）。具体的には、観測信号のマイク j とマイク j' とからの到来時間差 $q'_{jj'}$ を全てのマイクペアについて求め、それらを並べた縦ベクトルとマイクの座標系とから音声到来方向ベクトルを推定する。

【 0 0 0 7 】

10

20

30

40

50

各フレームごとに到来時間差 $q'_{jj'}$ を計算する手法として、非特許文献 4 にて開示されている G C C - P H A T と呼ばれる手法がある。この手法においては到来時間差 $q'_{jj'}(\tau)$ を次式に従い算出する。

【数 3】

$$q'_{jj'}(\tau) = \operatorname{argmax}_{q'} \sum_f \frac{x_j(f, \tau) \cdot x_{j'}^*(f, \tau)}{|x_j(f, \tau) \cdot x_{j'}^*(f, \tau)|} \cdot e^{j \cdot 2\pi f q'} \quad (3)$$

これをすべてのマイクペア $j j'$ について求めて、それらを並べた縦ベクトルを $vq'(\tau)$ とする。なお、すべてのマイクペアを用いる代わりに、ある基準マイクを決め、基準マイクとその他のマイクに関するすべてのペアを用いてもよい。音声到来方向ベクトル $vq(\tau)$ は、 $vq'(\tau)$ と音速 c とマイクの座標系 VD とから次式により推定する。

$$vq(\tau) = c \cdot VD^+ \cdot vq'(\tau) \quad (4)$$

ここで、 $^+$ は Moore-Penrose の疑似逆行列を表し、 vd_j がマイク j の座標を $[x, y, z]$ と並べたベクトルであるとき、 $VD = [vd_1 - vd_j, \dots, vd_M - vd_j]^T$ である。このように求めた音声到来方向ベクトル $vq(\tau)$ は、到来方向の水平角が θ 、仰角が φ とすると、次式のように表すことができる。

$$vq(\tau) = [\cos \theta \cdot \cos \varphi, \sin \theta \cdot \cos \varphi, \sin \varphi]^T \quad (5)$$

【0 0 0 8】

各フレームの各周波数成分ごとに到来時間差 $q'_{jj'}$ を計算する場合は、マイク j とマイク j' との到来時間差 $q'_{jj'}(f, \tau)$ を次式に従い算出する。

【数 4】

$$q'_{jj'}(f, \tau) = \frac{1}{2\pi f} \arg[x_j(f, \tau) \cdot x_{j'}^*(f, \tau)] \quad (6)$$

これをすべてのマイクペア $j j'$ について求めて（又は上記のように基準マイクに対して求めて）、それらを並べた縦ベクトルを $vq'(f, \tau)$ とし、式(4)と同様にして音声到来方向ベクトル $vq(f, \tau)$ を推定する。

【0 0 0 9】

なお、音声区間推定部 1 2 0 の処理と到来方向推定部 1 3 0 の処理とは並行して行ってもよいし、音声区間推定部 1 2 0 の処理により音声区間を推定した上で、その音声区間に該当するフレームに絞って到来方向推定部 1 3 0 の処理を行うこととしてもよい。

【0 0 1 0】

到来方向分類部 1 4 0 は、音声区間 P_s に該当する各フレームについて、音声到来方向（ベクトル $vq(\tau)$ ）又は $vq(f, \tau)$ が類似するものを各話者区間 P_k ($k = 1, \dots, N$) としてクラスタリングを行い、すべてのクラスタについて、クラスタのインデックス k とそのクラスタに属するすべてのフレームのインデックス τ との組を出力する（S 4）。

【数 5】

$$P_s = \sum_{k=1}^N P_k \quad (7)$$

【0 0 1 1】

クラスタリング手法としては、公知の k -m e a n s 法や階層的クラスタリングを用いてもよいし、オンラインクラスタリングを用いてもよい（非特許文献 5 参照）。このクラスタリング処理で分類されたクラスタ C_k が、そのクラスタを形成しているクラスタメンバ（ベクトル $vq(\tau)$ ）又は $vq(f, \tau)$ から求められるセントロイドで示される角度方向にいる話者 k に相当し、このクラスタメンバに該当する各フレーム τ が話者 k による話者区間 P_k を構成する。

【0 0 1 2】

なお、上記の説明では、到来方向推定部 1 3 0 はマイク間の到達時間差ベクトル $vq'(\tau)$

10

20

30

40

50

τ)又は $vq'(f, \tau)$ を推定した上で、更に音声到来方向ベクトル $vq(\tau)$ 又は $vq(f, \tau)$ を推定しているが、単に到達時間差ベクトルを推定するだけでも構わない。従って、この場合は図13に示すように、到来方向推定部130が到来時間差推定部131として構成され、到来方向分類部140が到来時間差分類部141として $vq(\tau)$ 又は $vq(f, \tau)$ の代わりに $vq'(\tau)$ 又は $vq'(f, \tau)$ を分類するように構成すればよい。

【特許文献1】特表2000-512108号公報

【非特許文献1】S.Araki, M.Fujimoto, K.Ishizuka, H.Sawada and S.Makino, "Speaker indexing and speech enhancement in real meetings/conversations," IEEE International Conference on Acoustics, Speech, and Signal Processing(ICASSP-2008), 2008, p.93-96

【非特許文献2】J.Sohn, N.S.Kim and W.Sung, "A Statistical Model-Based Voice Activity Detection," IEEE Signal Processing letters, 1999, vol.6, no.1, p.1-3

【非特許文献3】藤本、石塚、中谷、「複数の音声区間検出法の適応的統合の検討と考察」、電子情報通信学会 音声研究会、2007、SP2007-97、p.7-12

【非特許文献4】C.H.Knapp and G.C.Carter, "The generalized correlation method for estimation of time delay," IEEE Trans. Acoust. Speech and Signal Processing, 1976, vol.24, no.4, p.320-327

【非特許文献5】R.O.Duda, P.E.Hart and D.G.Stork, "Pattern Classification," 2nd edition, Wiley Interscience, 2000

【発明の開示】

【発明が解決しようとする課題】

【0013】

従来技術では、音声の到来方向情報のみにより話者識別を行っていたため、ある位置に居た話者が他の位置に移動してしまった場合に、同じ話者であるにもかかわらず新しい話者と識別したり、新しい話者であるにもかかわらず以前にその位置にいた別の話者として誤識別したりする問題があった。

本発明の目的は、音声の収録中に話者位置の移動が生じて、移動前と移動後において、同一話者には同一インデックスを付与することのできる、複数信号区間推定装置、複数信号区間推定方法、そのプログラムおよび記録媒体を提供することにある。

【課題を解決するための手段】

【0014】

本発明の複数信号区間推定装置は、複数のマイクによりそれぞれ収録された、複数の話者による発話音声が含まれる観測信号から、それぞれの話者の発話区間を推定するものであり、周波数領域変換部と音声区間推定部と到来方向推定部と到来方向分類部と話者同定部とを備える。

【0015】

周波数領域変換部は、観測信号を所定長のフレームに順次切り出し、当該フレームごとに周波数領域に変換する。

音声区間推定部は、周波数領域に変換された観測信号に基づき、各フレームが音声区間に該当するか否かを推定する。

到来方向推定部は、周波数領域に変換された観測信号に基づき、当該観測信号の到来方向を各フレームごとに推定する。

【0016】

到来方向分類部は、音声区間に該当すると推定された各フレームを、到来方向の類似性に基づき話者ごとのクラスタに分類する。

そして話者同定部は、所定の時刻までに同一クラスタに分類された各フレームの周波数領域に変換された観測信号に基づき、当該クラスタに係る話者のモデルをクラスタごとに作成し、当該所定の時刻以降の観測信号の話者を、各話者のモデルに基づき推定する。

【発明の効果】

【0017】

10

20

30

40

50

本発明の複数信号区間推定装置によれば、複数の話者に対する話者区間の推定にあたり、音声の到来方向の推定と分類に加え、話者の同一性の判定が可能となる。そのため、音声の収録中に話者位置の移動が生じて、移動前と移動後において、同一話者には同一のインデックスを付与することができる。

【発明を実施するための最良の形態】

【0018】

〔第1実施形態〕

図1（実線部分）に本発明の複数信号区間推定装置200の機能構成例を、図2（実線部分）にその処理フロー例を示す。複数信号区間推定装置200は、背景技術にて説明した周波数領域変換部110、音声区間推定部120、到来方向推定部130、及び到来方向分類部140と、話者同定部250とから構成される。また、話者同定部250の処理は図11に示したフローのS4に続いて行われる。従って、ここでは背景技術として説明した内容の説明は必要最小限とし、話者同定部250での処理に重点を置いて説明する。

図3（実線部分）に話者同定部250の機能構成例を示す。話者同定部250は、特徴抽出手段251とモデル学習手段252と尤度計算手段253とから構成される。

【0019】

話者同定部250の処理においては、観測信号の収録開始から所定の時刻 t_{train} までは話者の位置の移動が無かったと仮定し、その間に作成されたクラスタから、各話者のモデル M_k を作成することとする。そして、時刻 t_{train} 以降は話者の位置の移動があり得たと仮定し、時刻 t_{train} 以降のすべての音声セグメント（同一クラスタに分類された連続フレーム）について、その発話者が時刻 t_{train} 以前に発話したどの話者であるかを、観測信号の当初部分（収録開始から時刻 t_{train} まで）で作成した各話者のモデルに基づき判定する。このように各話者のモデルを観測信号の当初部分で作成することで、時刻 t_{train} 以降については、事前に話者のモデルを用意することなく話者の同定を行うことができる。なお、 t_{train} は同定の対象となる話者全員が少なくとも一度発話した時点以降の時刻に設定する。

【0020】

特徴抽出手段251は、M本のマイクにおける周波数領域の観測信号 $x_1(f, \tau) \sim x_M(f, \tau)$ の中から任意に選んだいずれか1本の観測信号 $x(f, \tau)$ の音声特徴量ベクトル $vf(\tau)$ を、各フレームごとに計算する（S5）。音声特徴量ベクトル $vf(\tau)$ としては、たとえば12次元のMFCC(Mel-Frequency Cepstrum Coefficient)を利用できる。また、自己相関法などで推定した基本周波数 $F0(\tau)$ を併用し、音声特徴量ベクトル $vf(\tau)$ の成分として含ませてもよい。

【0021】

モデル学習手段252は、到来方向分類部140にて同一クラスタ C_k （話者数Nのとき、 $k = 1, \dots, N$ ）に分類されたフレームのうち、観測信号の収録開始から所定の時刻 t_{train} までの各フレームに係る音声特徴量ベクトル $vf(\tau)$ を用いて、話者kのモデル、すなわちモデルパラメータ ϕ_k を作成して出力するとともに、所定の時刻 t_{train} までの各フレームのインデックス τ とそれらがそれぞれ属するクラスタに係る話者のインデックスkとの組を出力する（S6）。なお、同一話者のフレームが連続する場合は、各フレームのインデックスを出力する代わりに、連続フレームの始点と終点の時刻を出力してもよい。

【0022】

話者のモデルとしては、ここでは混合正規分布(GMM: Gaussian Mixture Model)を用いる場合を例示するが、他の話者同定や話者認識の方法（隠れマルコフモデルやベクトル量子化等）を用いてもよい。GMMのガウシアンを M_g とした時、モデル M_k のモデルパラメータを $\phi_k = (\text{平均 } \mu_{k,m}, \text{共分散行列 } \Sigma_{k,m}, \text{ガウシアン重み } w_{k,m})$ と置くと、GMMは次式のように表すことができる。

10

20

30

40

【数 6】

$$p(vf(\tau) | \phi_k) = \sum_{m=1}^{M_g} w_{k,m} \cdot p_{k,m}(vf(\tau)) \quad (8)$$

$$p_{k,m}(vf(\tau)) = \frac{1}{(2\pi)^{d/2} \cdot |\Sigma_{k,m}|^{1/2}} \cdot \exp\left\{-\frac{1}{2}(vf(\tau) - \mu_{k,m})^T \cdot (\Sigma_{k,m})^{-1} \cdot (vf(\tau) - \mu_{k,m})\right\} \quad (9)$$

ここで、 $p_{k,m}(vf(\tau))$ は話者 k の m 番目の多次元（次元数 d は音声特徴量ベクトルの次元と同じ）ガウシアンを表している。 M_g は例えば 10 とする。モデルパラメータ ϕ_k は、EM アルゴリズムなどを用いて、所定の時刻 t_{train} までのクラスター C_k に属する全てのフレームに基づき、次式によって求められる対数尤度 L が最大となる ϕ_k の値として計算することができる。

10

【数 7】

$$L = \sum_{\tau \in C_k, \tau < t_{train}} \log(p(vf(\tau) | \phi_k)) \quad (10)$$

ここで、EM アルゴリズムは、「汪他、”計算統計 I ～確率計算の新しい手法～」、岩波書店、2003、p158-162」等にて公知の技術である。

【0023】

なお、モデル学習部では、モデルパラメータ ϕ_k の推定精度を高める上で、各フレーム τ は互いに接続されていることが望ましい。そこで、接続されていない場合の処理方法の一例を説明する。図 4 (a) は観測信号の到来方向の時系列の例である。この例は、収録開始から時刻 t_{train} までの間に到来方向が $\theta_1 \rightarrow \theta_2 \rightarrow \theta_3 \rightarrow \theta_2 \rightarrow \theta_1$ の順に推移しており、つまり話者 1 → 話者 2 → 話者 3 → 話者 2 → 話者 1 の順に発話している場合である。このうち、話者 3 は短時間の隙間を挟んで計 3 回発話している。このように短時間（例えば 300ms 以下）の隙間があるような場合には、図 4 (b) に示すように音声区間が連続しているとみなしてモデルを学習するのが望ましい。また、話者 1 と話者 2 については、共に 1 回目の発話と 2 回目の発話との間が広がっている。このような場合には、図 4 (b) に示すように 1 回目の発話と 2 回目の発話が一体的にされたものとみなしてモデルを学習する。なお、モデル学習手段 252 が出力するインデックス τ は接続前の τ であることに注意が必要である。

20

30

【0024】

尤度計算手段 253 は、所定の時刻 t_{train} 以降に同一クラスターに分類された互いに接続されたフレーム（以下、「セグメント」という）の音声特徴量について、モデル学習手段 252 において作成した全ての話者のモデルに対する尤度を計算して、最大尤度をとるモデルに係る話者のインデックス k と当該セグメントに含まれる全てのフレームのインデックス τ とを出力する（S7）。なお、同一話者のフレームが連続する場合は、各フレームのインデックスを出力する代わりに、連続フレームの始点と終点の時刻を出力してもよい。

【0025】

話者のモデルとして GMM を用いた場合、各話者のモデルに当該セグメントに含まれる全てのフレーム τ の音声特徴量ベクトル $vf(\tau)$ を代入して、式 (10) により対数尤度を計算し、最も大きな対数尤度をとるモデルのインデックス k を当該セグメントの話者インデックスとして付与する。なお、話者の同定は必ずしもセグメントごとに行う必要はなく、フレームごとに行っても構わない。この場合、対数尤度の計算は式 (10) の Σ を外した式により行う。

40

【0026】

以上のように本発明においては、複数の話者に対する話者区間の推定にあたり、音声の到来方向の推定と分類に加え、話者の同定を行う。そのため、音声の収録中に話者位置の移動が生じて、移動前と移動後において、同一話者には同一のインデックスを付与することができる。

50

【 0 0 2 7 】

〔 第 2 実施形態 〕

第 1 実施形態においては、特徴抽出手段 2 5 1 における処理に際し、周波数領域変換部 1 1 0 から出力された周波数領域の観測信号 $x(f, \tau)$ をそのまま使用していた。しかし、実際の会議の場では複数の発話者がしばしば同時に発話するが、各フレームではいずれかの 1 名の話者の発話として識別する必要があり、その他の話者の発話は雑音成分となるため、同時発話されたフレーム τ における観測信号 $x(f, \tau)$ をそのまま使用すると、S/N 比の小ささにより特徴抽出を適切に行えずに話者モデルの推定精度が劣化する場合がある。そこで第 2 実施形態では、この S/N 比を向上させるための機能構成・処理方法を示す。

【 0 0 2 8 】

10

第 1 実施形態との機能構成上の相違は図 1 において、更に点線部分の構成、つまり音声強調部 2 6 0 が加わる点にあり、処理フロー上の相違は、図 2 において更に点線部分の処理が加わる点にある。

【 0 0 2 9 】

音声強調部 2 6 0 においては、それぞれの話者 k の発話信号成分を強調する。ここでは、複数のマイクにおける観測信号を用いた公知のビームフォーミング的手法（例えば、参考文献 1 参照）を用いてもよいし、1 本のマイクにおける観測信号に対して処理をする方法（例えば、Wiener Filter）による雑音除去的な手法を用いてもよい。

〔参考文献 1〕S. Araki, H. Sawada and S. Makino, "Blind Speech Separation in a Meeting Situation with Maximum SNR beamformers," proc. of ICASSP2007, 2007, vol.1, p.41-45

20

【 0 0 3 0 】

参考文献 1 の S/N 比最大化型ビームフォーマの場合には、周波数領域変換部 1 1 0 からの M 本のマイクにおける周波数領域の観測信号による観測信号ベクトル $\mathbf{v}_x(f, \tau) = [x_1(f, \tau), \dots, x_M(f, \tau)]^T$ と、到来方向分類部 1 4 0 からの各クラス C_k に属するフレーム τ の情報とから、各フレーム τ が属するクラス C_k に係る話者 k の発話信号成分を強調した周波数領域信号 $y_k(f, \tau)$ を生成し (S 8)、これを $x(f, \tau)$ の代わりに特徴抽出手段 2 5 1 での処理に用いる。

【 0 0 3 1 】

このように第 1 実施形態の構成に音声強調部 2 6 0 による処理を加えることで、特徴抽出手段 2 5 1 に入力する各話者 k の発話信号成分の S/N 比を向上することができ、話者モデルの推定精度を高めることができる。

30

【 0 0 3 2 】

〔 第 3 実施形態 〕

上記の実施形態では、モデルパラメータ ϕ_k を時刻 t_{train} までの観測信号により求めて、それを時刻 t_{train} 以降の話者同定処理に固定的に適用する。しかし、会話が収録される音響環境は通常、経時的に変化するものであり、求めたモデルパラメータ ϕ_k が経時的にその環境に相応しくなくなる場合がある。

【 0 0 3 3 】

第 3 実施形態はそのような事態を回避するための構成であり、処理フロー例を図 5 に示す。S 7 にて時刻 t_{train} 以降のセグメントに対して話者インデックス k を付与した後、そのセグメントに属する各フレームの音声特徴量ベクトル $\mathbf{v}_f(\tau)$ を、図 3 の一点鎖線に示すように尤度計算手段 2 5 4 からモデル学習手段 2 5 3 にフィードバックし、これらの音声特徴量ベクトル $\mathbf{v}_f(\tau)$ を用いて、式 (10) により改めて ϕ_k を計算してモデルパラメータを更新する (S 9)。更新は逐次行っても、所定の更新間隔を置いて行っても構わない。

40

【 0 0 3 4 】

このように構成することで、会話が収録される音響環境が経時的に変化しても、適切なモデルパラメータにより話者の同定処理を行うことができる。

【 0 0 3 5 】

〔 第 4 実施形態 〕

50

上記の各実施形態では、尤度計算手段 2 5 3 における話者の同定を、各話者のモデル M_k に同定対象セグメントに含まれる全てのフレーム τ の音声特徴量ベクトル $vf(\tau)$ を代入して対数尤度を計算し、対数尤度が最大となるモデルのインデックス k を当該セグメントの話者インデックスとするというルールの下で行う。しかし、このようなルールの下では、新たに参加した話者による発話があった場合においても、当初から参加している話者のモデルのいずれかが最大対数尤度をとることになるため、そのモデルの話者であると同定されてしまう。

【 0 0 3 6 】

第 4 実施形態はそのような事態を回避するための構成である。処理フロー例を図 6 に示す。尤度計算手段 2 5 3 において、所定の時刻 t_{train} 以降の各セグメントについて音声特徴量ベクトルを各話者のモデルに代入して対数尤度を計算し (S 7 - 1)、最大の対数尤度が所定の閾値より小さいか否かを判断し、閾値より大きい場合には、最大尤度をとるモデルに係る話者のインデックス k と当該セグメントに含まれる全てのフレームのインデックス τ とを出力し (S 7 - 2)、閾値より小さい場合には、新たな話者が参加したと判断して新たな話者インデックスを当該セグメントに付与するとともに、そのセグメントに属する各フレームの音声特徴量ベクトル $vf(\tau)$ を、図 3 の一点鎖線に示すように尤度計算手段 2 5 4 からモデル学習手段 2 5 3 にフィードバックし、これらの音声特徴量ベクトル $vf(\tau)$ を用いて、式 (10) により ϕ_k を計算して新たな話者のモデルパラメータとして追加する (S 1 0)。

【 0 0 3 7 】

このように構成することで、新たな話者が参加した場合においても、それを検知してその話者のモデルを生成することにより、以降、その話者についても同定処理を行うことができる。

【 0 0 3 8 】

〔 第 5 実施形態 〕

上記の各実施形態は、モデルパラメータを時刻 t_{train} までの観測信号により求めて、それを用いて時刻 t_{train} 以降の話者同定処理を行う構成である。しかし、発話が想定される複数の話者音声を予め入手できる場合には、それに基づき事前に各話者のモデルを準備しておき、この事前に準備したモデルを用いて話者同定処理を行うことが可能である。

第 5 実施形態はそのような場合の構成であり、話者同定部 2 5 0 を例えば図 7 のように構成することにより実現できる。上記の各実施形態との機能構成上の相違は、図 3 におけるモデル学習手段 2 5 2 が、予め準備した話者のモデルパラメータが記憶された話者モデル DB 2 6 4 に置き換わる点にある。

【 0 0 3 9 】

このように構成することで、モデルパラメータを学習により求める必要が無くなるため、音声の収録当初から尤度計算手段 2 5 3 において話者同定が可能になる。また、話者のモデルパラメータに話者の氏名情報を関連付けて DB に記憶させておくことで、話者インデックス k に方向情報に加え話者の氏名情報も持たせることができる。

上記の各実施形態の複数信号区間推定装置の構成をコンピュータによって実現する場合、各装置が有すべき機能の処理内容はプログラムによって記述される。そして、このプログラムをコンピュータで実行することにより、上記処理機能がコンピュータ上で実現される。

【 0 0 4 0 】

この処理内容を記述したプログラムは、コンピュータで読み取り可能な記録媒体に記録しておくことができる。コンピュータで読み取り可能な記録媒体としては、例えば、磁気記録装置、光ディスク、光磁気記録媒体、半導体メモリ等のようなものでもよいが、具体的には、例えば、磁気記録装置として、ハードディスク装置、フレキシブルディスク、磁気テープ等を、光ディスクとして、DVD (Digital Versatile Disc)、DVD-RAM (Random Access Memory)、CD-ROM (Compact Disc Read Only Memory)、CD-R (Recordable) / RW (ReWritable) 等を、光磁気記録媒体として、MO (Magneto-

10

20

30

40

50

Optical disc)等を、半導体メモリとしてEEPROM(Electronically Erasable and Programmable-Read Only Memory)等を用いることができる。

【0041】

また、このプログラムの流通は、例えば、そのプログラムを記録したDVD、CD-ROM等の可搬型記録媒体を販売、譲渡、貸与等することによって行う。さらに、このプログラムをサーバコンピュータの記憶装置に格納しておき、ネットワークを介して、サーバコンピュータから他のコンピュータにそのプログラムを転送することにより、このプログラムを流通させる構成としてもよい。

【0042】

また、上述した実施形態とは別の実行形態として、コンピュータが可搬型記録媒体から直接このプログラムを読み取り、そのプログラムに従った処理を実行することとしてもよく、さらに、このコンピュータにサーバコンピュータからプログラムが転送されるたびに、逐次、受け取ったプログラムに従った処理を実行することとしてもよい。また、サーバコンピュータから、このコンピュータへのプログラムの転送は行わず、その実行指示と結果取得のみによって処理機能を実現する、いわゆるASP(Application Service Provider)型のサービスによって、上述の処理を実行する構成としてもよい。なお、本形態におけるプログラムには、電子計算機による処理の用に供する情報であってプログラムに準ずるもの(コンピュータに対する直接の指令ではないがコンピュータの処理を規定する性質を有するデータ等)を含むものとする。

【0043】

また、この形態では、コンピュータ上で所定のプログラムを実行させることにより、本装置を構成することとしたが、これらの処理内容の少なくとも一部をハードウェア的に実現することとしてもよい。

また、上述の各種処理は、記載に従って時系列に実行されるのみならず、処理を実行する装置の処理能力あるいは必要に応じて並列的にあるいは個別に実行されてもよい。その他、本発明の趣旨を逸脱しない範囲で適宜変更が可能である。

【0044】

〔効果の確認〕

発明の効果を確認するため、図8で示すような3本のマイクを用いた測定環境において、4名参加による5分間の会議データについての話者区間推定実験を行った。会議においては、まず男女各2名の話者がそれぞれ男1、女1、男2、女2の位置に着席して始めに自己紹介をし、その後、各話者が順番に位置PPに移動して発言を行った。自己紹介は収録開始から120秒までの間に行われたものとし、 t_{train} を120秒として収録開始から120秒までの観測信号を話者同定モデルの作成に用い、120秒以降について話者同定を行った。なお、短時間フーリエ変換のフレーム長は64ms、フレームシフト長は32msとした。

【0045】

評価指標としては、diarization error rate(DER)を利用した。

【数8】

$$DER = \frac{\text{誤受理} \cdot \text{誤棄却} \cdot \text{話者誤りの時間長}}{\text{全音声区間長}} \times 100 (\%)$$

ここで、DERは誤棄却(missed speaker time: MST、誰かが話しているにもかかわらず話していないと判定した時間長)、誤受理(false alarm speaker time: FAT、誰も話していないにもかかわらず誰かが話していると判定した時間長)、話者誤り(speaker error time: SET、話者を誤って判定した時間長)の3つの誤検出を含む指標となっている。つまりこの指標においては、DER値が小さい方が話者区間推定の精度が高いことを示しており、特に本発明においては話者を正しく判定できているかが問題となるため、効果の程度はSETに顕著に現れるはずである。

【0046】

図 9 (a) に確認結果を示す。図 10 は結果を図解したものであり、(a) は正解を示したものの、(b) は従来の方法による推定結果、(c) は本発明の方法による推定結果である。なお、男 1、女 1、男 2、女 2 の到来方向はそれぞれ 100° 、 50° 、 -50° 、 -100° であり、位置 P P は -160° の到来方向にあり、また、男 1 が話者 1 に、女 1 が話者 2 に、男 2 が話者 3 に、女 2 が話者 4 にそれぞれ対応する。図 10 (b) からわかるように、従来の方法では位置 P P の話者を話者 1 ~ 4 以外の別の話者 5 と推定しており、図 9 (a) に示すとおり S E T が大きくなっている。これに対し、本発明の方法ではほぼ全ての時間区間で -160° 方向の話者の区別を図 10 (a) と同様にできており、図 9 (a) に示すとおり S E T が改善し、全体の性能である D E R 値も改善していることがわかる。

【0047】

10

また、10 組の話者組み合わせにおける会議シミュレーションを行った結果を図 9 (b) に示す。これは、音声信号と図 8 の測定環境で測定したインパルス応答とを用いて作成した会議シミュレーションデータを用いたものである。図 9 (b) においてシミュレーション 1 は各話者の音声間の重なりが無い場合であり、シミュレーション 2 は各話者の音声間の重なりがある場合の結果であるが、いずれの場合においても D E R、S E T に関し本発明の方法が従来方法より優れた結果を示すことがわかる。

【産業上の利用可能性】

【0048】

本発明は、複数話者の音声信号が混在している音響データから各話者の音声区間を推定する必要があるシステムや装置等に利用することができ、特に音声の収録中に話者位置の移動が生じる場合に有効である。

20

【図面の簡単な説明】

【0049】

【図 1】第 1、2 実施形態の複数信号区間推定装置の機能構成例を示す図

【図 2】第 1、2 実施形態の複数信号区間推定装置の処理フロー例を示す図

【図 3】第 1 ~ 4 実施形態の複数信号区間推定装置の話者同定部の機能構成例を示す図

【図 4】フレームが接続されていない場合に接続して処理をする方法を説明する図

【図 5】第 3 実施形態の複数信号区間推定装置の処理フロー例を示す図

【図 6】第 4 実施形態の複数信号区間推定装置の処理フロー例を示す図

【図 7】第 5 実施形態の複数信号区間推定装置の機能構成例を示す図

30

【図 8】効果の確認に用いた測定環境を示す図

【図 9】効果の確認結果を示す表

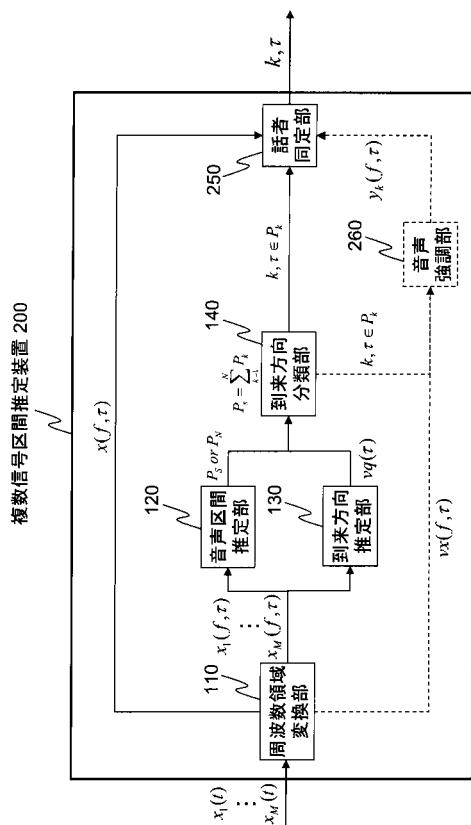
【図 10】効果の確認結果の根拠データを示す図

【図 11】従来技術の複数信号区間推定装置の機能構成例を示す図

【図 12】従来技術の複数信号区間推定装置の処理フロー例を示す図

【図 13】従来技術の複数信号区間推定装置の別の機能構成例を示す図

【図 1】



【図 3】

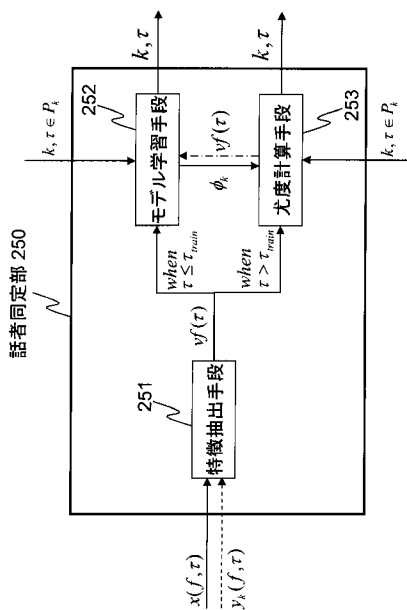


図3

(14)

JP 2010-54733 A 2010.3.11

【図 2】

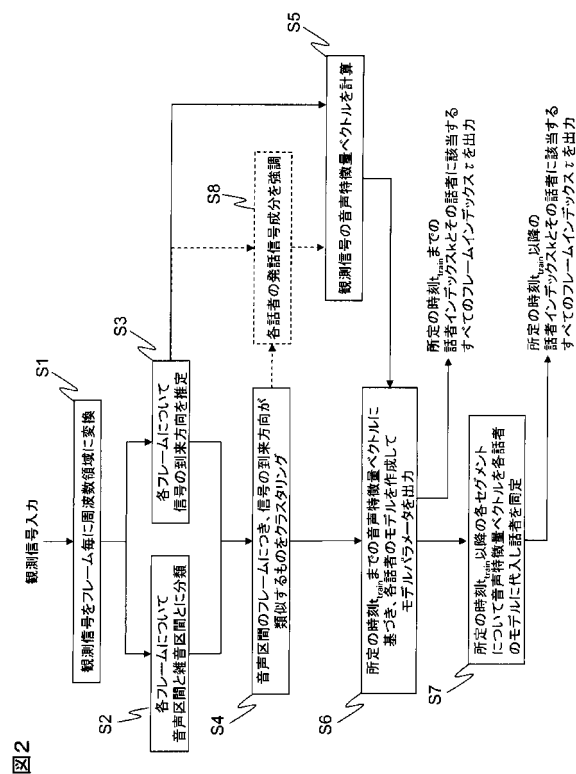


図2

【図 4】

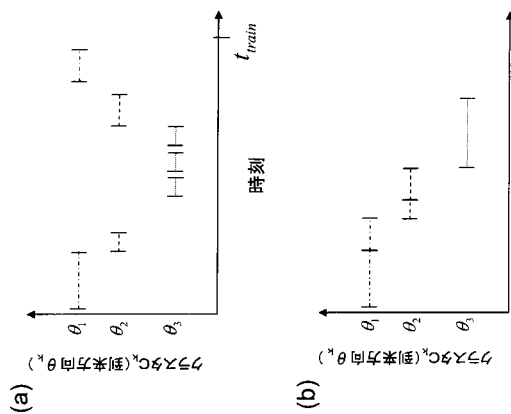
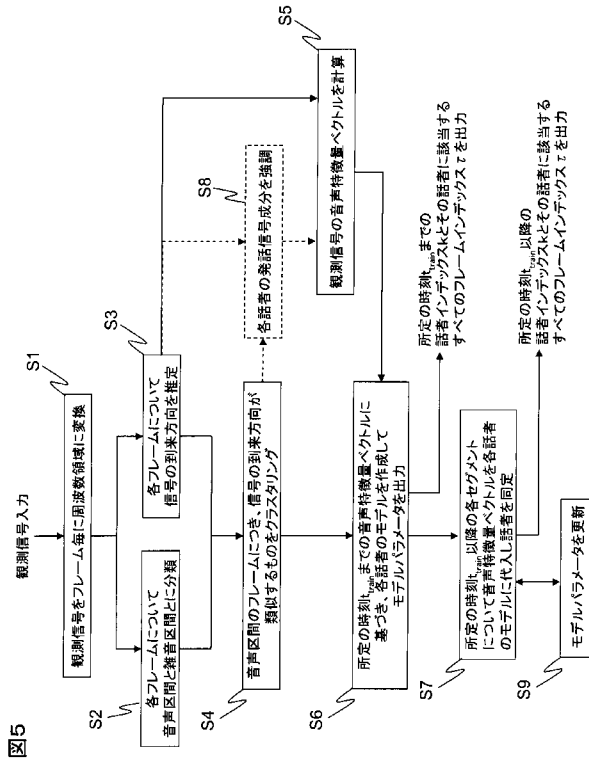
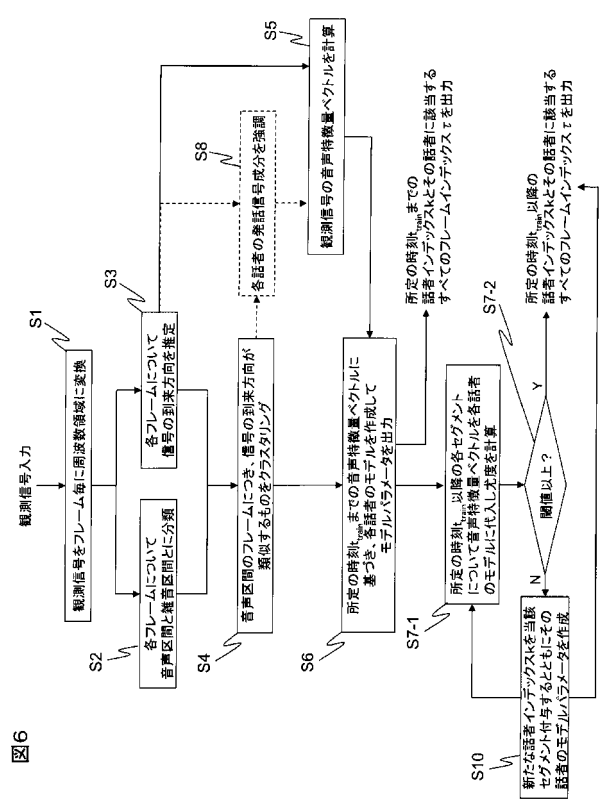


図4

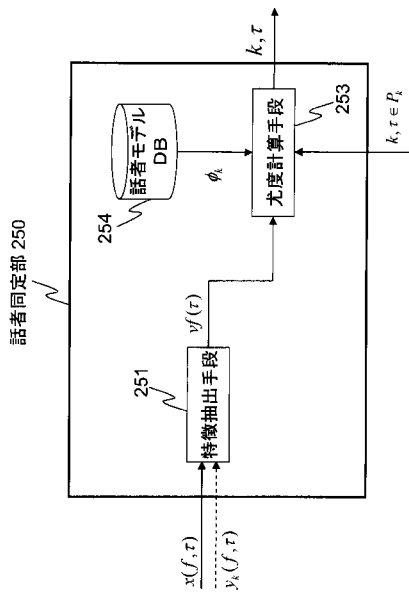
【図 5】



【図 6】



【図 7】



【図 8】

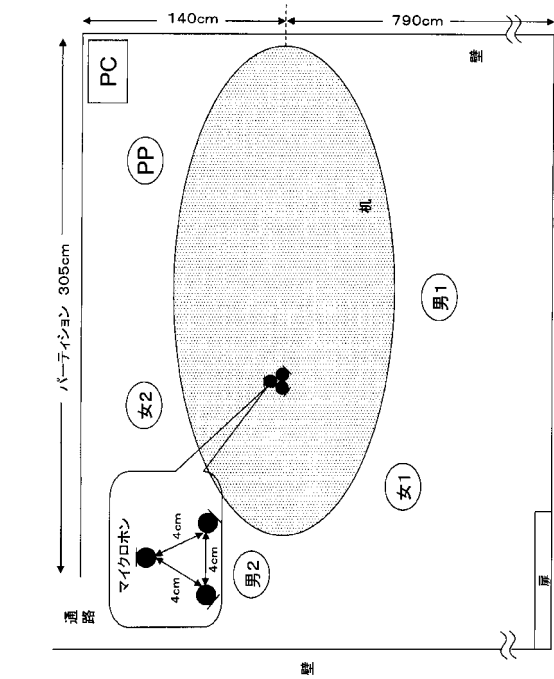


図7

図8

図9

(a)

	本発明(%)				従来方法(%)			
	DER	MST	FAT	SET	DER	MST	FAT	SET
	25.5	16.6	2.7	6.2	57.6	16.6	2.7	38.3
実測会議								

(b)

	本発明(%)				従来方法(%)			
	DER	MST	FAT	SET	DER	MST	FAT	SET
	8.7	2.1	0.1	6.4	34.8	2.2	0.1	32.5
シミュレーション1	18.3	10.1	1.3	6.9	39.9	10.6	0.2	29.1
シミュレーション2								

図11

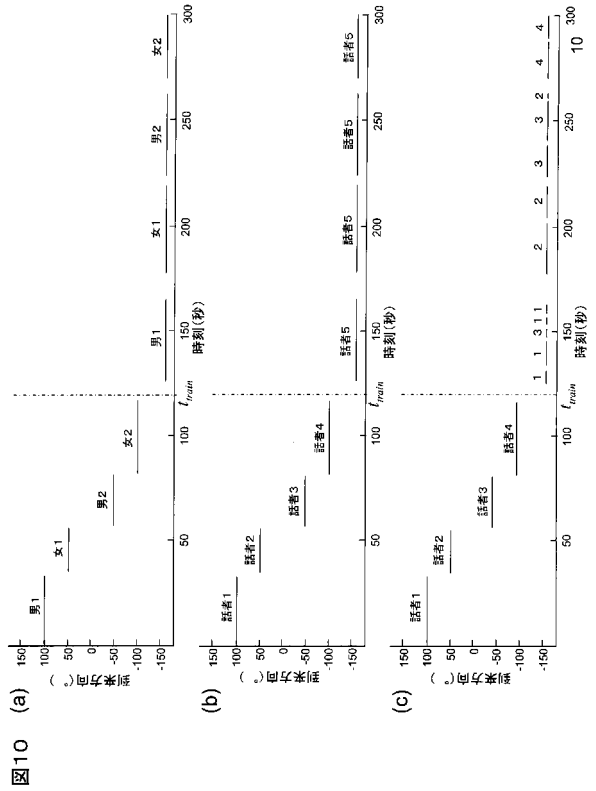
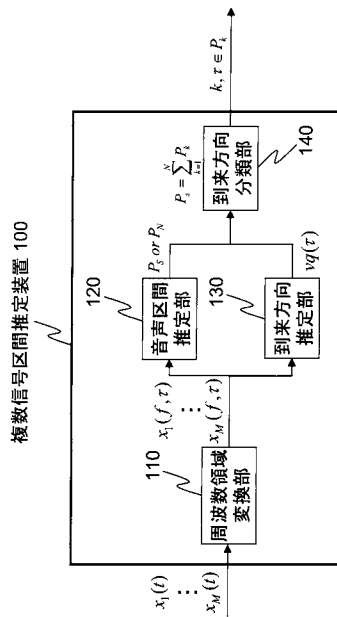


図12

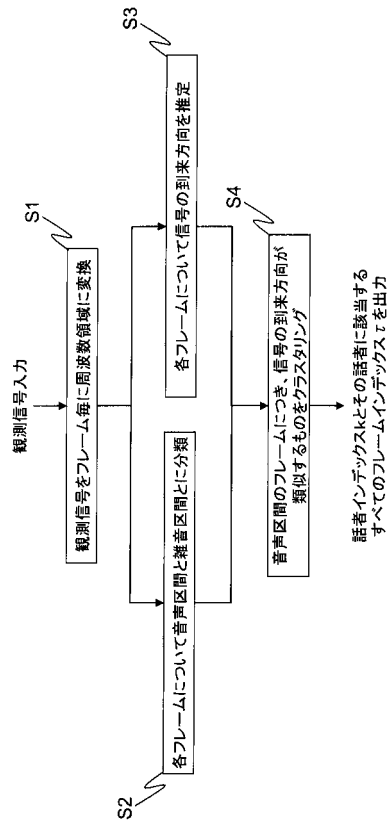
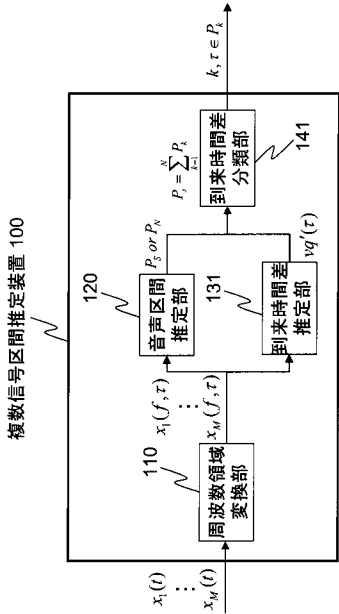


図13



フロントページの続き

- (72)発明者 藤本 雅清
東京都千代田区大手町二丁目3番1号 日本電信電話株式会社内
- (72)発明者 中谷 智広
東京都千代田区大手町二丁目3番1号 日本電信電話株式会社内
- (72)発明者 牧野 昭二
東京都千代田区大手町二丁目3番1号 日本電信電話株式会社内
- Fターム(参考) 5D015 AA03 DD03