

(19)日本国特許庁(JP)

(12)特許公報(B2)

(11)特許番号
特許第7548482号
(P7548482)

(45)発行日 令和6年9月10日(2024.9.10)

(24)登録日 令和6年9月2日(2024.9.2)

(51)国際特許分類

F I

G 1 0 L 25/51 (2013.01)

G 1 0 L 25/51

G 1 0 L 21/0272(2013.01)

G 1 0 L 21/0272 1 0 0 Z

請求項の数 14 (全23頁)

(21)出願番号	特願2023-528949(P2023-528949)	(73)特許権者	514187420
(86)(22)出願日	令和4年1月18日(2022.1.18)		テンセント・テクノロジー・(シェン
(65)公表番号	特表2023-549411(P2023-549411 A)		エン)・カンパニー・リミテッド
(43)公表日	令和5年11月24日(2023.11.24)		中華人民共和国 5 1 8 0 5 7 グアンド
(86)国際出願番号	PCT/CN2022/072460		ン, シェンジェン, ナンシャ・ディス
(87)国際公開番号	WO2022/156655		トリクト, ミッドウエスト・ディストリ
(87)国際公開日	令和4年7月28日(2022.7.28)		クト・オブ・ハイテックパーク ケジジ
審査請求日	令和5年5月16日(2023.5.16)		ョンギ・ロード テンセント・ビルディ
(31)優先権主張番号	202110083388.6	(74)代理人	100107766
(32)優先日	令和3年1月21日(2021.1.21)		弁理士 伊東 忠重
(33)優先権主張国・地域又は機関	中国(CN)	(74)代理人	100070150
			弁理士 伊東 忠彦
		(74)代理人	100135079
			弁理士 宮崎 修

最終頁に続く

(54)【発明の名称】 音声通話の制御方法、装置、コンピュータプログラム及び電子機器

(57)【特許請求の範囲】

【請求項1】

サーバが実行する、音声通話の制御方法であって、
混合された通話音声を取得するステップであって、前記混合された通話音声は、少なくとも1つの分岐音声を含む、ステップと、
前記通話音声に対して周波数領域変換を行い、前記通話音声の周波数領域における各周波数点に対応するエネルギー情報を決定するステップと、
ニューラルネットワークに基づいて各周波数点における前記エネルギー情報に対して分離処理を行い、各周波数点における対応する各分岐音声の該周波数点の全てのエネルギーにおけるエネルギー占有比率を決定するステップと、
各周波数点における各分岐音声の前記エネルギー占有比率に基づいて前記通話音声に含まれる分岐音声の数を決定するステップと、
前記分岐音声の数に基づいて、通話音声の制御方式を設定して前記音声通話を制御するステップと、を含む、方法。

【請求項2】

前記通話音声に対して周波数領域変換を行い、前記通話音声の周波数領域における各周波数点に対応するエネルギー情報を決定するステップは、
前記通話音声に対してフレーム分割処理を行い、少なくとも1つのフレームの音声情報を取得するステップと、
各フレームの前記音声情報に対して周波数領域変換を行い、周波数領域の音声エネルギー

ースペクトルを取得するステップと、

前記音声エネルギースペクトルに基づいて、前記通話音声の周波数領域における各周波数点に対応するエネルギー情報を決定するステップと、を含む、請求項 1 に記載の方法。

【請求項 3】

各フレームの前記音声情報に対して周波数領域変換を行い、周波数領域の音声エネルギースペクトルを取得するステップは、

時間領域の各フレームの前記音声情報に対してフーリエ変換を行い、各フレームの前記音声情報に対応する周波数領域の音声エネルギースペクトルを取得するステップ、を含む、請求項 2 に記載の方法。

【請求項 4】

前記音声エネルギースペクトルに基づいて、前記通話音声の周波数領域における各周波数点に対応するエネルギー情報を決定するステップは、

前記音声エネルギースペクトルにおける各周波数点に対応する振幅に対してモジュラス求め処理を行い、前記音声エネルギースペクトルに対応する振幅スペクトルを取得するステップと、

前記振幅スペクトルの二乗値を求め、前記二乗値に対して対数演算を行い、前記通話音声の周波数領域における各周波数点に対応するエネルギー情報を生成するステップと、を含む、請求項 2 又は 3 に記載の方法。

【請求項 5】

前記ニューラルネットワークは、長短期記憶ニューラルネットワークを含み、

ニューラルネットワークに基づいて各周波数点における前記エネルギー情報に対して分離処理を行い、各周波数点における対応する各分岐音声の該周波数点の全てのエネルギーにおけるエネルギー占有比率を決定するステップは、

前記エネルギー情報を予め設定された音声分離モデルに入力し、長短期記憶ニューラルネットワークに基づく畳み込み処理を行い、各周波数点における対応する分岐音声を決定するステップと、

各周波数点における対応する各分岐音声の該周波数点におけるエネルギー情報に基づいて、各周波数点における前記各分岐音声の該周波数点におけるエネルギー占有比率を決定するステップと、を含む、請求項 1 乃至 4 の何れかに記載の方法。

【請求項 6】

単一音声に対応する第 1 の音声サンプル、及び前記単一音声を含む混合音声に対応する第 2 の音声サンプルを取得するステップと、

前記第 1 の音声サンプルから第 1 の音声特徴を抽出し、前記第 2 の音声サンプルから第 2 の音声特徴を抽出するステップと、

前記第 2 の音声特徴を、長短期記憶人工ニューラルネットワークに基づいて構築された音声分離モデルに入力し、前記第 2 の音声特徴から分離された予測音声、及び前記予測音声の前記第 2 の音声サンプルにおける対応する予測エネルギー占有比率を決定するステップと、

前記第 1 の音声サンプルの前記第 2 の音声サンプルにおける実際エネルギー占有比率と前記予測エネルギー占有比率との比較結果に基づいて、前記音声分離モデルのパラメータを更新するステップと、をさらに含む、請求項 5 に記載の方法。

【請求項 7】

各周波数点における各分岐音声の前記エネルギー占有比率に基づいて前記通話音声に含まれる分岐音声の数を決定するステップは、

各分岐音声について、該分岐音声の各周波数点における対応するエネルギー占有比率に基づいて、該分岐音声の前記エネルギー占有比率の平均値を求めるステップと、

各分岐音声の前記平均値及び所定閾値に基づいて、前記通話音声に含まれる分岐音声の数を決定するステップと、を含む、請求項 5 又は 6 に記載の方法。

【請求項 8】

各分岐音声の前記平均値及び所定閾値に基づいて、前記通話音声に含まれる分岐音声の

10

20

30

40

50

数を決定するステップは、

各分岐音声の前記平均値と前記所定閾値との差の絶対値が差閾値よりも小さい場合、前記分岐音声の数が複数であると判定するステップと、

各分岐音声の前記平均値と前記所定閾値との差の絶対値が前記差閾値以上である場合、前記分岐音声の数が1つであると判定するステップと、を含む、請求項7に記載の方法。

【請求項9】

前記分岐音声の数が複数であると判定された場合、前記分岐音声の数に基づいて、通話音声の制御方式を設定して前記音声通話を制御するステップは、

設定された音声抽出方式に基づいて、主要発言者の音声を抽出するステップ、を含む、請求項1乃至8の何れかに記載の方法。

【請求項10】

前記分岐音声の数が複数であると判定された場合、設定された音声抽出方式に基づいて、主要発言者の音声を抽出するステップは、

各周波数点における複数の分岐音声のそれぞれに対応するエネルギー占有比率に基づいて、前記エネルギー占有比率のうちの最大値に対応する分岐音声を前記主要発言者の音声として認識するステップと、

前記エネルギー情報から前記主要発言者の音声に対応する周波数情報を決定するステップと、

前記周波数情報に基づいて前記通話音声から前記主要発言者の音声を抽出するステップと、を含む、請求項9に記載の方法。

【請求項11】

前記分岐音声の数は、1つ又は少なくとも2つを含み、

前記分岐音声の数に基づいて、通話音声の制御方式を設定して前記音声通話を制御するステップは、

前記分岐音声の数が1つである場合、設定されたシングルトークのエコー処理方式に基づいて、前記分岐音声のエコー音声を認識し、前記エコー音声に対してシングルトークエコー除去を行うステップと、

前記分岐音声の数が少なくとも2つである場合、設定されたダブルトークのエコー処理方式に基づいて、前記分岐音声に対応するエコー音声をそれぞれ認識し、前記エコー音声に対してダブルトークエコー除去を行うステップと、を含む、請求項1乃至10の何れかに記載の方法。

【請求項12】

音声通話の制御装置であって、

混合された通話音声を取得する取得部であって、前記混合された通話音声は、少なくとも1つの分岐音声を含む、取得部と、

前記通話音声に対して周波数領域変換を行い、前記通話音声の周波数領域における各周波数点に対応するエネルギー情報を決定する変換部と、

ニューラルネットワークに基づいて各周波数点における前記エネルギー情報に対して分離処理を行い、各周波数点における対応する各分岐音声の該周波数点の全てのエネルギーにおけるエネルギー占有比率を決定する分離部と、

各周波数点における各分岐音声の前記エネルギー占有比率に基づいて前記通話音声に含まれる分岐音声の数を決定する数決定部と、

前記分岐音声の数に基づいて、通話音声の制御方式を設定して前記音声通話を制御する制御部と、を含む、装置。

【請求項13】

プロセッサにより実行される際に、請求項1乃至11の何れかに記載の音声通話の制御方法を実現する、コンピュータプログラム。

【請求項14】

1つ又は複数のプロセッサと、

1つ又は複数のプログラムを記憶する記憶装置と、を含む電子機器であって、

前記１つ又は複数のプログラムは、前記１つ又は複数のプロセッサにより実行される際に、前記１つ又は複数のプロセッサに請求項１乃至１１の何れかに記載の音声通話の制御方法を実現させる、電子機器。

【発明の詳細な説明】

【技術分野】

【０００１】

本発明は、２０２１年１月２１日に出願した出願番号が２０２１１００８３３８８．６であり、発明の名称が「音声通話の制御方法、装置、コンピュータ読み取り可能な記憶媒体及び電子機器」である中国特許出願に基づく優先権を主張し、その全ての内容を参照により本発明に援用する。

【０００２】

本発明の実施例は、コンピュータ技術の分野に関し、特に音声通話の制御方法、装置、コンピュータ読み取り可能な記憶媒体及び電子機器に関する。

【背景技術】

【０００３】

多くの音声通話のシナリオでは、その後の音声制御のために、発言者の数や音色などを判別する必要がある。関連技術では、多数のラベル付きの音声セグメントに基づいて、発言者シナリオの検出システムを訓練する。ここで、各セグメントのラベルは発言者の数であり、テストの際に１つの音声セグメントを与え、システムは現在の発言者の数を予測する。このような処理方式は、音声検出に長い遅延をもたらす、特にリアルタイムの通信シナリオにおいて、音声認識の効率を大幅に低下し、リアルタイムの音声制御効果に影響を与えてしまう。

【発明の概要】

【０００４】

本発明の実施例は、少なくともある程度で音声人数の検出精度を保証することができると共に、音声人数の識別効率及び音声通話の制御効率を向上させることができる、音声通話の制御方法、装置、コンピュータ読み取り可能な記憶媒体及び電子機器を提供する。

【０００５】

本発明の他の特徴及び利点は、以下の詳細な説明によって明らかになり、又は本発明の実施により部分的に明らかになるであろう。

【０００６】

本発明の実施例の１つの態様では、音声通話の制御方法であって、混合された通話音声を取得するステップであって、前記混合された通話音声は、少なくとも１つの分岐音声を含む、ステップと、前記通話音声に対して周波数領域変換を行い、前記通話音声の周波数領域における各周波数点に対応するエネルギー情報を決定するステップと、ニューラルネットワークに基づいて各周波数点における前記エネルギー情報に対して分離処理を行い、各周波数点における対応する各分岐音声の該周波数点の全てのエネルギーにおけるエネルギー占有比率を決定するステップと、各周波数点における各分岐音声の前記エネルギー占有比率に基づいて前記通話音声に含まれる分岐音声の数を決定するステップと、前記分岐音声の数に基づいて、通話音声の制御方式を設定して前記音声通話を制御するステップと、を含む、方法を提供する。

【０００７】

本発明の実施例の１つの態様では、音声通話の制御装置であって、混合された通話音声を取得する取得部であって、前記混合された通話音声は、少なくとも１つの分岐音声を含む、取得部と、前記通話音声に対して周波数領域変換を行い、前記通話音声の周波数領域における各周波数点に対応するエネルギー情報を決定する変換部と、ニューラルネットワークに基づいて各周波数点における前記エネルギー情報に対して分離処理を行い、各周波数点における対応する各分岐音声の該周波数点の全てのエネルギーにおけるエネルギー占有比率を決定する分離部と、各周波数点における各分岐音声の前記エネルギー占有比率に基づいて前記通話音声に含まれる分岐音声の数を決定する数決定部と、前記分岐音声の数

10

20

30

40

50

に基づいて、通話音声の制御方式を設定して前記音声通話を制御する制御部と、を含む、装置を提供する。

【0008】

本発明の実施例の1つの態様では、コンピュータプログラムが記憶されたコンピュータ読み取り可能な記憶媒体であって、前記コンピュータプログラムは、プロセッサにより実行される際に、本発明の実施例に記載の音声通話の制御方法を実現する、記憶媒体を提供する。

【0009】

本発明の実施例の1つの態様では、1つ又は複数のプロセッサと、1つ又は複数のプログラムを記憶する記憶装置と、を含む電子機器であって、前記1つ又は複数のプログラムは、前記1つ又は複数のプロセッサにより実行される際に、前記1つ又は複数のプロセッサに本発明の実施例に記載の音声通話の制御方法を実現させる、電子機器を提供する。

10

【0010】

本発明の実施例の1つの態様では、コンピュータ読み取り可能な記憶媒体に記憶されたコンピュータ命令を含むコンピュータプログラム製品又はコンピュータプログラムであって、コンピュータ装置のプロセッサは、コンピュータ読み取り可能な記憶媒体からコンピュータ命令を読み取り、該コンピュータ命令を実行することによって、前記コンピュータ装置に本発明の実施例に記載の音声通話の制御方法を実行させる、コンピュータプログラム製品又はコンピュータプログラムを提供する。

【0011】

なお、上記の一般的な説明及び後述の詳細な説明は、単なる例示的なもの及び解釈的なものであり、本発明を限定するものではない。

20

【図面の簡単な説明】

【0012】

【図1】本発明の実施例の技術を適用可能な例示的なシステムアーキテクチャを示す概略図である。

【図2】本発明の幾つかの実施例に係る会議システムを概略的に示す概略図である。

【図3】本発明の幾つかの実施例に係る音声通話の制御方法を概略的に示すフローチャートである。

【図4】本発明の幾つかの実施例に係る音声分離の流れを概略的に示す概略図である。

30

【図5】本発明の幾つかの実施例に係る音声抽出を概略的に示す概略図である。

【図6】本発明の幾つかの実施例に係る会議音声抽出を概略的に示す概略図である。

【図7】本発明の幾つかの実施例に係るエコー除去の適用シナリオを概略的に示す図である。

【図8】本発明の幾つかの実施例に係るエコー除去を概略的に示す概略図である。

【図9】本発明の幾つかの実施例に係るエネルギー情報の抽出を概略的に示すフローチャートである。

【図10】本発明の幾つかの実施例に係るエネルギー情報の抽出を概略的に示す概略図である。

【図11】本発明の幾つかの実施例に係る分離モデルの訓練を概略的に示すフローチャートである。

40

【図12】本発明の幾つかの実施例に係る主要発言者の設定のインターフェースを概略的に示す図である。

【図13】本発明の幾つかの実施例に係る通信リソースの割り当てのインターフェース図を概略的に示す図である。

【図14】本発明の幾つかの実施例に係る音声通話の制御装置を概略的に示すブロック図である。

【図15】本発明の実施例の電子機器を実現可能なコンピュータシステムの構成を示す概略図である。

【発明を実施するための形態】

50

【0013】

クラウドコンピューティング (cloud computing) は計算モードの一種であり、計算タスクを大量のコンピュータにより構成されたリソースプールに分布することによって、各種の応用システムが需要に応じて計算力、記憶空間及び情報サービスを取得できる。リソースを提供するネットワークは「クラウド」と呼ばれる。「クラウド」内のリソースは、ユーザの目には無限に拡張可能であり、いつでも利用可能であり、オンデマンドで、いつでも拡張可能であり、従量課金される。クラウドコンピューティングの基礎能力の提供者として、クラウドコンピューティングリソースプール (クラウドプラットフォームと略称し、一般的にIaaS (Infrastructure as a Service: サービスとしてのインフラストラクチャ) プラットフォーム) と呼ばれ、リソースプールの中に多種類の仮想リソースを構成し、外部の顧客が選択して使用することができる。クラウドコンピューティングリソースプールは、主に計算機器 (仮想化機器であり、オペレーティングシステムを含む)、ストレージ機器、ネットワーク機器を含む。論理的な機能別では、IaaS (Infrastructure as a Service: サービスとしてのインフラストラクチャ) 層にPaaS (Platform as a Service: サービスとしてのプラットフォーム) 層を配備し、PaaS層にSaaS (Software as a Service: サービスとしてのソフトウェア) 層を配備してもよいし、IaaSにSaaSを直接配備してもよい。PaaSは、例えばデータベース、ウェブコンテナなどのソフトウェアを実行するプラットフォームである。SaaSは、例えばWebポータルやSMSグループプロバイダなどの様々なサービスソフトウェアである。一般的には、SaaSやPaaSはIaaSと比べて上位層である。

10

20

【0014】

クラウドコールセンター (Cloud Call Center) は、クラウドコンピューティング技術に基づいて構築されたコールセンターシステムであり、企業はいかなるソフトウェア、ハードウェアシステムを購入する必要がなく、人員、場所などの基本条件を備えるだけで、迅速に自分のコールセンターを所有することができ、ソフトウェアとハードウェアのプラットフォーム、通信リソース、日常のメンテナンスとサービスはサーバ業者から提供される。建設周期が短く、投入が少なく、リスクが低く、配備が柔軟であり、システム容量の拡張性が強く、運営メンテナンスコストが低いなどの特徴がある。電話マーケティングセンター、顧客サービスセンターにかかわらず、企業は必要に応じてサービスをレンタルするだけで、機能が全面的で、安定的で、信頼性があり、座席が全国各地に分布でき、全国コールが接続できるコールセンターシステムを確立することができる。

30

【0015】

本発明の実施例では、クラウドコールセンターの方式でセンターシステムをコールすることができ、同時に前記システムに音声通話の制御方法を埋め込み、コール中の音声制御を実現し、さらにクラウドコールセンターをよりスマート化させ、クラウドコールセンターの信頼性と安全性を向上させる。

【0016】

クラウド会議は、クラウドコンピューティング技術に基づく効率的で、便利で、低コストの会議形式である。ユーザは、インターネットインターフェースを通じて、簡単に使いやすい操作を行うだけで、迅速かつ効率的に世界各地のチームと顧客と同時に音声、データファイルとビデオを共有することができ、会議中のデータの転送、処理などの複雑な技術はクラウド会議サービス業者がユーザの操作を助ける。現在、国内クラウド会議は、主にSaaS (Software as a Service: サービスとしてのソフトウェア) モードを主体とするサービス内容に集中し、電話、ネットワーク、ビデオなどのサービス形式を含み、クラウドコンピューティングに基づくビデオ会議はクラウド会議と呼ばれる。クラウド会議時代に、データの転送、処理、保存は全てビデオ会議メーカーのコンピュータリソースによって処理され、ユーザは、高価なハードウェアとインストールの面倒なソフトウェアを購入する必要がなく、ブラウザを開き、対応するインターフェースを登録するだけで、効率的な遠隔会議を行うことができる。クラウド会議システムは、マル

40

50

チサーバの動的クラスタ配備をサポートし、かつ複数の高性能サーバを提供し、会議の安定性、安全性、可用性を大幅に向上させる。近年、ビデオ会議は、コミュニケーション効率を大幅に高め、コミュニケーションコストを持続的に低下させ、内部管理レベルのアップグレードをもたらすため、多くのユーザに歓迎され、すでに政府、軍隊、交通、運輸、金融、運営者、教育、企業などの各分野に広く応用されている。勿論、ビデオ会議がクラウドコンピューティングを利用した後、利便性、快速性、使いやすさの面で更に強い魅力があり、きっとビデオ会議応用の新しいクライマックスの到来につながる。

【0017】

クラウド会議の応用シナリオにおいて、本発明の実施例は、音声通話に基づく制御方法をクラウド会議に応用することができ、クラウド会議の過程における音声通話をよりはっきりさせ、音声通信過程をよりインテリジェント化させ、さらに会議の効率を向上させる。

10

【0018】

クラウドソーシャル (Cloud Social) は、Internet of Things (IoT)、クラウドコンピューティングとモバイルインターネットの相互作用応用の仮想ソーシャルアプリケーションモデルの1つであり、有名な「リソース共有関係マップ」を創立することを目的とし、さらにネットソーシャルを展開する。クラウドソーシャルの主要な特徴は、大量の社会リソースを統一的に整合と評価し、1つのリソース有効プールを構成してユーザにオンデマンドでサービスを提供することである。共有に参加するユーザが増えれば増えるほど、生み出せる利用価値は大きくなる。

【0019】

人工知能 (Artificial Intelligence: AI) は、人間の知能をシミュレーションし、拡張し、環境を感知し、知識を取得し、知識を使用して最適な結果を得るために、デジタルコンピュータ又はデジタルコンピュータによって制御される機械を利用する理論、方法、技術及び応用システムである。言い換えれば、人工知能は、計算機科学の総合技術であり、知能の本質を理解し、人間の知能と同様の方法で反応する新しい知能機械を生産しようとする。人工知能は、各種の知能機械の設計原理と実現方法を研究し、機械に感知、推理と決定の機能を持たせる。人工知能技術は、1つの総合的な学科であり、領域が広く、ハードウェアレベルの技術もソフトウェアレベルの技術もある。人工知能基礎技術は一般的に例えばセンサー、専用人工知能チップ、クラウドコンピューティング、分散型ストレージ、ビッグデータ処理技術、操作／インタラクティブシステム、メカトロニクスなどの技術を含む。人工知能ソフトウェア技術は、主にコンピュータビジョン技術、音声処理技術、自然言語処理技術、及び機械学習／深層学習などの幾つかの方面を含む。

20

【0020】

音声技術 (Speech Technology) のキーテクノロジーには、自動音声認識 (ASR) と音声合成 (TTS)、声紋認識がある。コンピュータが聞くことができ、見ることができ、話すことができ、感じるができるようにすることは、未来のヒューマンコンピュータインタラクションの発展方向であり、その中で音声は、未来の最も有望なヒューマンコンピュータインタラクション方式の1つとなっている。機械学習 (Machine Learning: ML) は1つの多分野の交差学科であり、確率論、統計学、近似論、凸分析、算法複雑度理論などの多学科に関わる。コンピュータがどのように人間の学習行為を模擬或いは実現するかを専念に研究し、新しい知識或いは技能を取得し、既存の知識構造を再組織し、絶えず自身の性能を改善させる。機械学習は、人工知能の核心であり、コンピュータに知能を持たせる根本的なルートであり、その応用は人工知能の各領域に及んでいる。機械学習とディープラーニングは、通常、人工ニューラルネットワーク、信頼ネットワーク、強化学習、遷移学習、帰納学習式教育学習などの技術を含む。

30

40

【0021】

人工知能技術の研究と進歩に伴い、人工知能技術は多くの領域で研究と応用を展開し、例えば一般的なスマートホーム、スマートウェアデバイス、バーチャルアシスタント、スマートスピーカー、スマートマーケティング、無人運転、自動運転、ドローン、ロボット

50

、知能医療、知能カスタマーサービスなど、技術の発展に伴い、人工知能技術は更に多くの領域で応用され、ますます重要な価値を発揮すると信じている。

【0022】

関連技術では、多数のラベル付きの音声セグメントに基づいて、発言者シナリオの検出システムを訓練する。ここで、各セグメントのラベルは発言者の数であり、テストの際に1つの音声セグメントを与え、システムは現在の発言者の数を予測する。しかし、このスキームは、検出過程において、現在の発言者数を判定するために多くのコンテキスト情報を必要とし、例えば、発言者の数を決定するために、比較的に長い時間の音声セグメントをデータ基礎として取り込む必要がある。このような処理方式は、音声検出に長い遅延をもたらし、特にリアルタイムの通信シナリオにおいて、音声認識の効率を大幅に低下し、リアルタイムの音声制御効果に影響を与えてしまう。

10

【0023】

本発明の実施例は、人工知能の音声技術及び機械学習等の技術に関するものであり、これらの技術により、本発明の実施例に係る音声通話の制御方法をより正確にすることができる。具体的には、以下の実施例を参照しながら説明する。

【0024】

図1は、本発明の実施例の技術を適用可能な例示的なシステムアーキテクチャを示す概略図である。

【0025】

図1に示すように、システムアーキテクチャは、端末装置（図1に示すように、スマートフォン101、タブレットコンピュータ102及びポータブルコンピュータ103のうちの1つ以上であってもよいし、デスクトップコンピュータなどであってもよい）、ネットワーク104、及びサーバ105を含むことができる。ネットワーク104は、端末装置とサーバ105との間の通信リンクを提供するための媒体である。ネットワーク104は、有線通信リンク、無線通信リンクなどの様々な接続タイプを含むことができる。

20

【0026】

なお、図1における端末装置、ネットワーク及びサーバの数は、単に例示的なものである。実装要件に応じて、任意の数の端末装置、ネットワーク、及びサーバを有することができる。例えば、サーバ105は、複数のサーバからなるサーバクラスタ等であってもよい。

30

【0027】

なお、本実施例の各端末装置は、異なる通話用クラスタを対象とすることができ、通話用クラスタ内の参加者数は、1人、2人、又はそれ以上などであってもよい。例えば、ポータブルコンピュータ103を対象とする通話クラスタには複数の参加者が含まれてもよく、タブレットコンピュータ102を対象とする通話クラスタには他の参加者が含まれてもよく、ユーザはスマートフォン101を介して会議に参加してもよい。

【0028】

一例として、会議の進行中に、複数のユーザ又は1人のユーザが端末装置を使用して会議通話を行うことができる。同時に、サーバ105は、ネットワーク104を介して端末装置との通話音声を取得し、通話音声に対して周波数領域変換を行い、通話音声の周波数領域に対応するエネルギー情報を決定することができる。ニューラルネットワークに基づいてエネルギー情報に対して分離処理を行い、通話音声に含まれる各分岐音声の通話音声におけるエネルギー占有比率を決定する。エネルギー占有比率に基づいて通話音声に含まれる分岐音声の数を決定する。分岐音声の数に基づいて、通話音声制御方式を設定して音声通話を制御する。

40

【0029】

上記のスキームは、通話進行過程においてリアルタイムに通話音声を取得し、通話音声に対して周波数領域変換を行って通話音声の周波数領域に対応するエネルギー情報を決定し、その後、ニューラルネットワークに基づいてエネルギー情報に対して分離処理を行い、通話音声に含まれる各分岐音声の通話音声におけるエネルギー占有比率を決定し、エネ

50

ルギー占有比率に基づいて通話音声に含まれる分岐音声の数を決定し、最後に、分岐音声の数に基づいて、通話音声制御方式を設定することによって音声通話を制御する。これによって、音声通話過程における音声人数の即時検出、及び音声通話のリアルタイム制御を実現し、音声人数の検出精度を保証すると共に、音声人数の認識効率及び音声通話の制御効率を向上させることができる。

【0030】

これに加えて、図2に示すように、本実施例では、1つの通話クラスタのみを対象として処理することも可能であり、該通話クラスタは、1つ、2つ、又は複数の参加者を含む。上述した音声通話の制御方法によれば、通話クラスタ内のリアルタイムの発言者の数を検出し、通話中の音声品質を保証すると共に、対応する通話制御を行うことで、通話効率

10

【0031】

なお、本発明の実施例に係る音声通話の制御方法は、一般にサーバ105によって実行されるため、音声通話の制御装置は、一般にサーバ105に設けられる。しかしながら、本発明の他の実施例では、端末装置は、本発明の実施例に係る音声通話の制御スキームを実行するために、サーバと同様の機能を有することもできる。

【0032】

なお、本実施例におけるサーバは、独立した物理サーバであってもよいし、複数の物理サーバからなるサーバクラスタや分散型システムであってもよいし、クラウドコンピューティングサービスを提供するクラウドサーバであってもよい。端末としては、スマートフォン、タブレットコンピュータ、ノートパソコン、デスクトップパソコン、スマートスピーカー、スマートウォッチ等が挙げられるが、これらに限定されない。端末及びサーバは、有線又は無線通信により直接又は間接に接続することができるが、本発明はここに限定されない。

20

【0033】

以下は、本発明の実施例の技術的解決策の実施の詳細を説明する。

【0034】

図3は、本発明の幾つかの実施例に係る音声通話の制御方法を概略的に示すフローチャートである。該音声通話の制御方法はサーバにより実行されてもよく、該サーバは図1に示すサーバであってもよい。図3に示すように、該音声通話の制御方法は、少なくともステップS110～ステップS150を含み、以下に詳細に説明する。

30

【0035】

ステップS110において、混合された通話音声を取得し、該混合された通話音声は、少なくとも1つの分岐音声を含む。

【0036】

本発明の幾つかの実施例では、通話クラスタが通話を行う間に、混合された通話音声を取得してもよい。本実施例の通話音声の長さは、制限されず、リアルタイムで取得された1フレームの通話音声であってもよいし、時間長が1秒間又は1分間の通話音声であってもよい。

【0037】

例えば、該通話は、リアルタイム通信会議のシナリオであってもよい。該リアルタイム通信会議の間に、リアルタイムで通話音声を収集し、収集された通話音声に基づいて対応する認識処理を行い、生成された識別結果に基づいてその後の制御を行い、通話音声のリアルタイム制御の効果を達成する。

40

【0038】

ステップS120において、該通話音声に対して周波数領域変換を行い、該通話音声の各周波数点に対応するエネルギー情報を決定する。

【0039】

本発明の1つの実施例では、本実施例で取得される通話音声は、時間領域の通話音声であり、時間を独立変数とし、音量を従属変数とする音声信号である。本実施例では、通話

50

音声を取得した後、通話音声に対して周波数領域変換を行い、時間領域の音声信号を周波数領域の音声信号に変換して、通話音声の周波数領域でのエネルギー情報を表す。

【0040】

図4に示すように、音声に基づいて分離された発言者シーンの分類フレームワークは、「信号前処理」段階において、本発明の実施例では、オーディオ信号を取得し、オーディオ信号を前処理することによって音響特徴を抽出し、オーディオ信号に対応する対数エネルギースペクトルを、通話音声の周波数領域における各周波数点に対応するエネルギー情報として生成する。

【0041】

具体的には、本実施例におけるエネルギー情報は、通話音声の各周波数点に対応するエネルギー値、エネルギースペクトル等の情報を含んでもよい。本実施例では、エネルギー情報によって各周波数点のエネルギーなどの属性を評価し、エネルギー情報に基づいて各周波数点における対応する分岐音声を区別することができる。

【0042】

ステップS130において、ニューラルネットワークに基づいて各周波数点におけるエネルギー情報に対して分離処理を行い、各周波数点における対応する各分岐音声の該周波数点の全てのエネルギーにおけるエネルギー占有比率を決定する。

【0043】

本発明の1つの実施例では、各周波数点におけるエネルギー情報を取得した後、ニューラルネットワークに基づいてエネルギー情報に対して分離処理を行う。即ち、既に訓練された分離モデルにエネルギー情報を入力して、各分岐音声の該周波数点におけるエネルギー占有比率を取得し、該エネルギー占有比率は、音声分離に基づく周波数点係数とも呼ばれる。例えば、図4において、2つの分岐音声を一例にすると、1つの周波数点における2つの分岐音声の周波数点係数、即ち、周波数点係数 P_A 及び周波数点係数 P_B を取得することができる。なお、図4では、単なる2つの分岐音声を一例として、1つの周波数点における2つの周波数点係数を示しており、それぞれ発言者Aと発言者Bに対応しているが、実際には、1つの周波数点で得られる周波数点係数は、同時に話している人の数に関連し、2つに限らない。また、エネルギー値が0でない各周波数点では、何れも音声分離に基づく周波数点係数が得られる。

【0044】

具体的には、本実施例における周波数点係数は、1つの周波数点における発言者に対応するエネルギーの、該周波数点の全てのエネルギー情報に占める割合を表すものである。本実施例における音声分離の考え方は、周波数領域の各周波数点係数の方式に基づいており、ある発言者がある周波数点におけるエネルギー占有比率は、混合信号において予測された周波数点係数の大きさに正比例する。周波数点係数（ P ）の計算方式は、1つの周波数点におけるある発言者の音声エネルギー値（ E ）と該周波数点における混合発言者の音声エネルギー値との比である。二人（AとB）を仮定し、以下の式に従って1つの周波数点における発言者Aの周波数点係数を計算する。

【0045】

【数1】

$$P_A = \frac{E_A}{E_A + E_B}$$

以上の式の計算により1つの周波数点における周波数点係数 P_A 及び周波数点係数 P_B が得られた後、 P_A が P_B よりも大きい場合、該周波数点は発言者Aが主に行うものであり、 P_A が P_B よりも小さい場合、該周波数点は発言者Bが主に行うものであ

【0046】

上記の方法は、特に複数人が同時に発話するシーンにおいて、エネルギー情報を分離することによって、各周波数点における各分岐音声に対応するエネルギー占有比率を決定し

、エネルギー占有比率に基づいて各分岐音声の分布状況を決定し、音声数の認識の正確性とリアルタイム性を向上させることができる。

【0047】

ステップS140において、各周波数点における各分岐音声のエネルギー占有比率に基づいて通話音声に含まれる分岐音声の数を決定する。

【0048】

本発明の1つの実施例では、通話音声中の各周波数点における各分岐音声のエネルギー占有比率を決定した後、本実施例では、エネルギー占有比率に基づいて、通話音声に含まれる分岐音声の数を平均化することにより決定する。

【0049】

本発明の1つの実施例では、発言者Aの音声の各フレーム内の各周波数点に対応するエネルギー占有比率について、各フレーム内の各周波数点のエネルギー占有比率を平均化し、1フレーム時間内の安定したエネルギーの平均値を取得し、そして、所定の閾値に基づいて現在のエネルギーの平均値が一人の発言に対応するか、それとも複数の人の発言に対応するかを判定し、最後に、現在フレームの発言者数情報を出力する。例えば、各フレームの現在の発言者数に対応する離散的な0（一人の発言）又は1（複数の人の発言）を出力することができる。

【0050】

本発明の1つの実施例では、収集された通話音声は多数のフレームからなり、1フレーム中に複数の周波数点が存在する。例えば、周波数点の個数は、フーリエ変換されたポイント数であってもよく、1フレーム中の周波数点の個数を f とし、 F_i はある発言者のその中の i 番目の周波数点に対応するエネルギー占有比率、即ち周波数点係数であり、平均値を求めることによって、該発言者の該フレームに対応するエネルギー占有比率の平均値は、以下ようになる。

【0051】

【数2】

$$\frac{1}{f} \sum_{i=0}^{i=f-1} F_i。$$

そして、各分岐音声に対する平均値と閾値とを比較することによって、同時に発話する人の数を決定することができる。例えば、2人の場合、 $P_A + P_B = 1$ であるため、 P_A 及び P_B のうちの何れかと閾値とを比較すればよい。例えば、 P_A を一例にすると、実際の音声人数判定では、その値が0又は1である場合、現在発話しているエネルギー（ P_A 又は P_B ）が音声エネルギー全体を占めていることを示しているため、ある1人だけが発話している、即ちB又はAが発話していることを意味する。この場合、発言者数は1である。その値が0.5である場合、2人とも同時に発話しており、且つその時の発話エネルギーの大きさが同一であることを意味する。この場合、発言者数は2である。計算されたエネルギー占有比率の平均値と所定の閾値とを比較することで、現在の分岐音声の数を決定することができる。実際の応用では、応用シナリオに応じて閾値の具体的な値を設定することができる。

【0052】

本実施例では、上記の閾値検出方式により現在の分岐音声の数を決定し、リアルタイムにフレームレベルの非常に短い時間内で多発言者シナリオを判断し、リアルタイムに音声ストリームを処理することができる。また、多対多のラベルを使用し、音声情報を十分に活用し、シナリオ検出の正確率を向上させることができる。

【0053】

ステップS150において、決定された分岐音声の数に基づいて、通話音声の制御方式を設定して前記音声通話を制御する。

【0054】

本発明の1つの実施例では、分岐音声の数を決定した後、現在の分岐音声の数に基づいて現在の通話状況を決定し、さらに設定された通話音声制御方式により音声通話を制御してもよい。これによって、音声通話のリアルタイム制御を実現し、音声制御の精度性とリアルタイム性を向上させることができる。

【0055】

例えば、図5に示すように、分岐音声の数が複数である場合、分岐音声の数に基づいて、背景発言者をフィルタリングにより除去し、主要な発言者のオーディオストリームのみを抽出する。この際、フロントエンドで現在発言者の数を検出する必要があり、現在の発言者の数が1よりも大きい場合、主要発言者抽出をオンにし、現在の発言者数が1であることを検出した場合、音声への損傷を回避するために、音声抽出をオフにする。

10

【0056】

本実施例では、音声を抽出する過程では、各周波数点における複数の分岐音声のそれぞれに対応するエネルギー占有比率に基づいて、エネルギー占有比率のうちの最大値に対応する分岐音声を主要発言者の音声として認識し、エネルギー情報から主要発言者の音声に対応する周波数情報を決定し、周波数情報に基づいて通話音声から主要発言者の音声を抽出する。

【0057】

図6に示すように、複数のユーザが発話しているシーンでは、上述した周波数検出方式により、そのうちの主要発言者、例えば図6におけるユーザ4を特定してもよい。また、音声通話を明瞭にするように、主要発言者の音声を抽出し、或いは他のユーザの音声をフィルタリングにより除去する。

20

【0058】

以上のように、複数の人が話すシーンでその中の一人の主要発言者の発話音声を抽出することができ、通話中の音声をよりはっきりさせ、通話品質と効果を向上させることができる。

【0059】

図7に示す音声を外部に出力する場合、通話者の一方は、他方から戻ってきた音声から自分のエコーを聞こえ、通話品質が低下するという問題がある。

【0060】

この問題を避けるために、図8に示すように、通信相手側と自分側とが交互に話しているシングルトークのシナリオの場合、即ち、分岐音声の数が1つである場合、設定されたシングルトークのエコー処理方式に基づいて、分岐音声のエコー音声を認識し、エコー音声に対してシングルトークエコー除去を行う。

30

【0061】

通信相手側と自分側とが同時に発話するダブルトークのシナリオの場合、即ち、分岐音声の数が少なくとも2つである場合、設定されたダブルトークのエコー処理方式に基づいて、分岐音声に対応するエコー音声をそれぞれ認識し、エコー音声に対してダブルトークエコー除去を行う。通信システムでは、できるだけ自分側の信号がエコー除去過程において最大限に保留できることを保証する。

40

【0062】

上記のスキームは、通話進行過程においてリアルタイムに通話音声を取得し、通話音声に対して周波数領域変換を行って通話音声の周波数領域に対応するエネルギー情報を決定し、その後、ニューラルネットワークに基づいてエネルギー情報に対して分離処理を行い、通話音声に含まれる各分岐音声の通話音声におけるエネルギー占有比率を決定し、エネルギー占有比率に基づいて通話音声に含まれる分岐音声の数を決定し、最後に、分岐音声の数に基づいて、通話音声制御方式を設定することによって音声通話を制御する。これによって、音声通話過程における音声人数の即時検出、及び音声通話のリアルタイム制御を実現し、音声人数の検出精度を保証すると共に、音声人数の認識効率及び音声通話の制御効率を向上させることができる。

50

【0063】

本発明の1つの実施例では、図9に示すように、ステップS120における通話音声に対して周波数領域変換を行い、通話音声の周波数領域における各周波数点に対応するエネルギー情報を決定するプロセスは、ステップS1210～ステップS1230を含み、以下のように詳細に説明する。

【0064】

ステップS1210において、通話音声に対してフレーム分割処理を行い、少なくとも1つのフレームの音声情報を取得する。

【0065】

ステップS1220において、各フレームの音声情報に対して周波数領域変換を行い、周波数領域の音声エネルギースペクトルを取得する。

10

【0066】

ステップS1230において、音声エネルギースペクトルに基づいて、通話音声の周波数領域における各周波数点に対応するエネルギー情報を決定する。

【0067】

本発明の1つの実施例では、各フレームの前記音声情報に対して周波数領域変換を行い、周波数領域の音声エネルギースペクトルを取得するステップは、時間領域の各フレームの音声情報に対してフーリエ変換（他の時間領域を周波数領域に変換する方式を含む）を行い、各フレームの音声情報に対応する周波数領域の音声エネルギースペクトルを取得する。

20

【0068】

本実施例のステップS1230における音声エネルギースペクトルに基づいて、通話音声の周波数領域における各周波数点に対応するエネルギー情報を決定するステップは、音声エネルギースペクトルにおける各周波数点に対応する振幅に対してモジュラス求め処理を行い、音声エネルギースペクトルに対応する振幅スペクトルを取得するステップと、振幅スペクトルの二乗値を求め、二乗値に対して対数演算を行い、通話音声の周波数領域における各周波数点に対応するエネルギー情報を生成するステップと、を含む。

【0069】

図10に示すように、本発明の1つの実施例では、収集された時間領域音声に対してフレーム分割処理、ウィンドウ追加処理を行い、その後、フレーム毎にそれぞれN点フーリエ変換を行い、フーリエ変換して得られたN個の周波数点のフーリエ変換係数を求め、それに対してモジュラスを求めて周波数領域の振幅スペクトルを取得し、得られた振幅スペクトルに対して二乗を求めて対数エネルギースペクトルを取得し、音声のN個の周波数点におけるエネルギー情報を取得する。

30

【0070】

本発明の1つの実施例では、ニューラルネットワークは、長短期記憶ニューラルネットワークを含む。ステップS130におけるニューラルネットワークに基づいて各周波数点におけるエネルギー情報に対して分離処理を行い、各周波数点における対応する各分岐音声の該周波数点の全てのエネルギーにおけるエネルギー占有比率を決定するステップは、エネルギー情報を予め設定された音声分離モデルに入力し、長短期記憶ニューラルネットワークに基づく畳み込み処理を行い、各周波数点における、通話音声に含まれる各分岐音声の該周波数点に対応するエネルギー占有比率を決定するステップを含む。

40

【0071】

本発明の1つの実施例では、該方法は、図11に示すように、音声分離モデルを訓練するプロセスにおいて、以下のステップをさらに含む。

【0072】

ステップS1110において、単一音声に対応する第1の音声サンプル、及び単一音声を含む混合音声に対応する第2の音声サンプルを取得する。

【0073】

ステップS1120において、第1の音声サンプルから第1の音声特徴を抽出し、第2

50

の音声サンプルから第2の音声特徴を抽出する。

【0074】

ステップS1130において、第2の音声特徴を、長短期記憶人工ニューラルネットワークに基づいて構築された音声分離モデルに入力し、第2の音声特徴から分離された予測音声、及び予測音声の第2の音声サンプルにおける対応する予測エネルギー占有比率を決定する。

【0075】

ステップS1140において、第1の音声サンプルの第2の音声サンプルにおける実際エネルギー占有比率と予測エネルギー占有比率との比較結果に基づいて、音声分離モデルのパラメータを更新する。

【0076】

本発明の1つの実施例では、まず、訓練用のデータセットを構築して2つの音声ライブラリ、即ち、一人発話用のコーパスと複数人発話用のコーパスを取得する。ここで、一人発話用のコーパスには、単一の音声に対応する第1の音声サンプルが含まれ、複数人発話用のコーパスは、それぞれランダムに複数の一人発話用のセグメントを抽出して重ね合わせ、その後、それぞれこの2つのデータベースに対して信号前処理により音声中の対数エネルギースペクトル特徴を抽出し、さらに分離モデルを経てそれぞれ音声セグメントの周波数係数を取得し、さらに後処理により分離後の音声を取得する。本実施例では、抽出された対数エネルギースペクトルの特徴を入力とし、この特徴を2層の長短期記憶ネットワーク(Long Short-Term Memory: LSTM)及び1層の出力層からなる分離モデルに入力して、周波数点係数を取得する。

【0077】

なお、本実施例でLSTMネットワークを用いる理由は、現時点の入力を考慮することだけでなく、ネットワークに前の内容の記憶機能を付与するためである。同時に、本実施例のネットワーク構造における追加した入力ゲート、出力ゲート、忘却ゲート、細胞状態ユニットは、LSTMのタイミングモデリング能力を著しく向上させ、より多くの情報を記憶することができ、データの長時間依存性を効果的に把握することができる。

【0078】

モデル全体の訓練では、訓練セットにマルチ発言者とシングル発言者が含まれており、複数のシングル発言者音声を用いて加算してマルチ発言者音声を取得し、シングル発言者はミュートとの混合と見なすことができる。ここで、混合音声から音声特徴aを抽出し、対応するきれいな音声から音声特徴bを抽出し、特徴ベクトルaを訓練入力とし、特徴ベクトルbを訓練目標とし、モデル出力は周波数点係数m、nとする。後処理により分離された音声を取得し、分離された音声ときれいな音声との誤差でLSTMモデルを訓練し、LSTMモデルにおけるパラメータを調整することによって、最終の分離モデルを取得し、得られた分離モデルを更に精確で完全にすることができる。

【0079】

上記の訓練プロセスで得られた音声分離モデルは、リアルタイムに多発言者シナリオをフレームレベルで短い時間で判断し、音声ストリームをリアルタイムに処理することができる。また、フレームにおける各周波数点に対応するラベルに基づいて、複数のラベルに複数の周波数点を対応付けることで、音声情報を十分に活用し、シナリオ検出の正確率を向上させることができる。

【0080】

さらに、本実施例では、ステップS1130における第2の音声特徴を、長短期記憶人工ニューラルネットワークに基づいて構築された音声分離モデルに入力し、第2の音声特徴から分離された予測音声を決定するプロセスにおいて、得られた周波数係数に混合音声の周波数スペクトルを乗算し、さらに逆フーリエ変換を経て、混合信号の位相を結合し、分離された音声信号を取得することができる。

【0081】

ステップS140では、各周波数点における各分岐音声のエネルギー占有比率に基づい

10

20

30

40

50

て通話音声に含まれる分岐音声の数を決定するステップは、各分岐音声について、該分岐音声の各周波数点における対応するエネルギー占有比率に基づいて、該分岐音声のエネルギー占有比率の平均値を求めるステップと、各分岐音声の平均値及び所定閾値に基づいて、通話音声に含まれる分岐音声の数を決定するステップと、を含む。

【0082】

本発明の1つの実施例では、収集された通話音声は多数のフレームからなり、1フレーム中に複数の周波数点が存在する。1フレーム中の周波数点の個数を f とし、 F_i はある発言者のその中の i 番目の周波数点に対応するエネルギー占有比率、即ち周波数点係数であり、平均値を求めることによって、該発言者の該フレームに対応するエネルギー占有比率の平均値は、以下のようになる。

【0083】

【数3】

$$\frac{1}{f} \sum_{i=0}^{i=f-1} F_i。$$

本発明の1つの実施例では、各分岐音声の平均値及び所定閾値に基づいて、通話音声に含まれる分岐音声の数を決定するステップは、各分岐音声の平均値と所定閾値との差の絶対値が差閾値よりも小さい場合、分岐音声の数が複数であると判定するステップと、各分岐音声の平均値と所定閾値との差の絶対値が差閾値以上である場合、分岐音声の数が1つであると判定するステップと、を含む。

【0084】

具体的には、本実施例では、2人が同時に発話することを一例にすると、平均値が0に近いほど、又は1に近いほど、1人発話の確率が大きくなり、0.5に近いほど、2人同時発話の確率が大きくなる。閾値の決定は、具体的なタスクに応じて決定される。例えば、主要発言者抽出のアルゴリズムのように、1人発話の場合にアルゴリズムが音声への損傷を避けるために、1人発言者の誤検出率が低いと判断する必要がある、その場合には閾値は0又は1に近い値に設定してもよい。

【0085】

図12に示すように、実際の会議適用シナリオにおいて、現在の発言者数が複数の人であると検出された場合、検出された発言者又は参加者をインターフェースに表示し、ユーザがトリガした主要発言者設定の指示に応じて、そのうちの何れか1人又は複数の人を主要発言者として設定し、残りの人の音声をフィルタリングにより除去して、会議の通話品質を保証することができる。

【0086】

図13に示すように、会話制御サーバは、複数の音声会話が同時に行われている場合、発言者数の多い会議により多くの通信リソースを割り当て、通話の品質を保証することができる。

【0087】

以下は、本発明の上記の実施例における音声通話の制御方法を実行可能な本発明の装置の実施例を説明する。なお、該装置は、コンピュータ装置内で実行される1つのコンピュータプログラム（プログラムコードを含む）であってもよく、例えば、該装置は、アプリケーションソフトウェアである。該装置は、本発明の実施例に係る方法における対応するステップを実行するために使用されてもよい。本発明の装置の実施例で開示されていない詳細について、本発明の上記の音声通話の制御方法の実施例を参照してもよい。

【0088】

図14は、本発明の幾つかの実施例に係る音声通話の制御装置を概略的に示すブロック図である。

【0089】

図14に示すように、本発明の1つの実施例に係る音声通話の制御装置1400は、以下の各部を含む。取得部1410は、混合された通話音声を取得し、混合された通話音声は、少なくとも1つの分岐音声を含む。変換部1420は、通話音声に対して周波数領域変換を行い、通話音声の周波数領域における各周波数点に対応するエネルギー情報を決定する。分離部1430は、ニューラルネットワークに基づいて各周波数点におけるエネルギー情報に対して分離処理を行い、各周波数点における対応する各分岐音声の該周波数点の全てのエネルギーにおけるエネルギー占有比率を決定する。数決定部1440は、各周波数点における各分岐音声のエネルギー占有比率に基づいて通話音声に含まれる分岐音声の数を決定する。制御部1450は、分岐音声の数に基づいて、通話音声の制御方式を設定して音声通話を制御する。

10

【0090】

本発明の幾つかの実施例では、上記の実施例をベースとして、変換部1420は、通話音声に対してフレーム分割処理を行い、少なくとも1つのフレームの音声情報を取得するフレーム分割部と、各フレームの音声情報に対して周波数領域変換を行い、周波数領域の音声エネルギースペクトルを取得する周波数領域変換部と、音声エネルギースペクトルに基づいて、通話音声の周波数領域における各周波数点に対応するエネルギー情報を決定するエネルギー決定部と、を含む。

【0091】

本発明の幾つかの実施例では、上記の実施例をベースとして、周波数領域変換部は、時間領域の各フレームの音声情報に対してフーリエ変換を行い、各フレームの前記音声情報に対応する周波数領域の音声エネルギースペクトルを取得する。

20

【0092】

本発明の幾つかの実施例では、上記の実施例をベースとして、エネルギー決定部は、音声エネルギースペクトルにおける各周波数点に対応する振幅に対してモジュラス求め処理を行い、音声エネルギースペクトルに対応する振幅スペクトルを取得し、振幅スペクトルの二乗値を求め、二乗値に対して対数演算を行い、通話音声の周波数領域における各周波数点に対応するエネルギー情報を生成する。

【0093】

本発明の幾つかの実施例では、上記の実施例をベースとして、ニューラルネットワークは、長短期記憶ニューラルネットワークを含む。分離部1430は、エネルギー情報を予め設定された音声分離モデルに入力し、長短期記憶ニューラルネットワークに基づく畳み込み処理を行い、各周波数点における対応する分岐音声を決定し、各周波数点における対応する各分岐音声の該周波数点におけるエネルギー情報に基づいて、各周波数点における各分岐音声の該周波数点におけるエネルギー占有比率を決定する。

30

【0094】

本発明の幾つかの実施例では、上記の実施例をベースとして、音声通話の制御装置1400は、更新部をさらに含む。更新部は、単一音声に対応する第1の音声サンプル、及び単一音声を含む混合音声に対応する第2の音声サンプルを取得し、第1の音声サンプルから第1の音声特徴を抽出し、第2の音声サンプルから第2の音声特徴を抽出し、第2の音声特徴を、長短期記憶人工ニューラルネットワークに基づいて構築された音声分離モデルに入力し、第2の音声特徴から分離された予測音声、及び予測音声の第2の音声サンプルにおける対応する予測エネルギー占有比率を決定し、第1の音声サンプルの第2の音声サンプルにおける実際エネルギー占有比率と予測エネルギー占有比率との比較結果に基づいて、音声分離モデルのパラメータを更新する。

40

【0095】

本発明の幾つかの実施例では、上記の実施例をベースとして、数決定部1440は、各分岐音声について、該分岐音声の各周波数点における対応するエネルギー占有比率に基づいて、該分岐音声のエネルギー占有比率の平均値を求める平均部と、各分岐音声の平均値及び所定閾値に基づいて、通話音声に含まれる分岐音声の数を音声数決定部と、を含む。

【0096】

50

本発明の幾つかの実施例では、上記の実施例をベースとして、音声数決定部は、各分岐音声の平均値と所定閾値との差の絶対値が差閾値よりも小さい場合、分岐音声の数が複数であると判定する第1の数判定部と、各分岐音声の平均値と所定閾値との差の絶対値が差閾値以上である場合、分岐音声の数が1つであると判定する第2の数判定部と、を含む。

【0097】

本発明の幾つかの実施例では、制御部1450は、設定された音声抽出方式に基づいて、主要発言者の音声抽出する抽出部、を含む。

【0098】

本発明の幾つかの実施例では、上記の実施例をベースとして、抽出部は、各周波数点における複数の分岐音声のそれぞれに対応するエネルギー占有比率に基づいて、エネルギー占有比率のうちの最大値に対応する分岐音声を主要発言者の音声として認識し、エネルギー情報から主要発言者の音声に対応する周波数情報を決定し、周波数情報に基づいて通話音声から主要発言者の音声抽出する。

【0099】

本発明の幾つかの実施例では、上記の実施例をベースとして、分岐音声の数は、1つ又は少なくとも2つを含む。制御部1450は、分岐音声の数が1つである場合、設定されたシングルトークのエコー処理方式に基づいて、分岐音声のエコー音声を認識し、エコー音声に対してシングルトークエコー除去を行い、分岐音声の数が少なくとも2つである場合、設定されたダブルトークのエコー処理方式に基づいて、分岐音声に対応するエコー音声をそれぞれ認識し、エコー音声に対してダブルトークエコー除去を行う。

【0100】

図15は、本発明の実施例の電子機器を実現可能なコンピュータシステムの構成を示す概略図である。

【0101】

なお、図15に示す電子機器のコンピュータシステム1500は一例に過ぎず、本発明の実施例の機能や使用範囲に何ら制限を加えるべきものではない。

【0102】

図15に示すように、コンピュータシステム1500は、読み出し専用メモリ1502 (Read-Only Memory: ROM) に記憶されたプログラムや、記憶部1508からランダムアクセスメモリ1503 (Random Access Memory: RAM) にロードされたプログラムに応じて、各種の適宜の動作や処理を実行可能な中央処理部1501 (Central Processing Unit: CPU) を有する。ランダムアクセスメモリ1503には、システム動作に必要な各種プログラムやデータも記憶されている。中央処理部1501、読み出し専用メモリ1502及びランダムアクセスメモリ1503は、バス1504を介して互いに接続されている。入力/出力インターフェース1505 (Input/Outputインターフェース、即ちI/Oインターフェース) もバス1504に接続される。

【0103】

入力部1506 (キーボード、マウスなどを含む)、出力部1507 (ディスプレイ、例えば陰極線管 (Cathode Ray Tube: CRT)、液晶ディスプレイ (Liquid Crystal Display: LCD) など、及びスピーカなどを含む)、記憶部1508 (例えばハードディスクなどを含む)、通信部1509 (例えばネットワークのインタフェースカード、例えばLANカード、モデムなどを含む) は、入力/出力インターフェース1505に接続されている。通信部1509は、ネットワーク、例えばインターネットを介して通信処理を実行する。必要に応じて、ドライバ1510は、入力/出力インターフェース1505に接続されてもよい。取り外し可能な媒体1511は、例えば磁気ディスク、光ディスク、光磁気ディスク、半導体メモリなどであり、必要に応じてドライバ1510にセットアップされて、その中から読みだされたコンピュータプログラムは必要に応じて記憶部1508にインストールされている。

【0104】

10

20

30

40

50

特に、本発明の実施例によれば、様々な方法のフローチャートに記載されたプロセスは、コンピュータソフトウェアプログラムとして実装することができる。例えば、本発明の実施例は、コンピュータ読み取り可能な媒体に記憶されたコンピュータプログラムを含むコンピュータプログラム製品を含み、コンピュータプログラムは、フローチャートに示された方法を実行するためのプログラムコードを含む。このような実施例では、コンピュータプログラムは、通信部 1509 を介してネットワークにダウンロードされてインストールされ、且つ／或いは、取り外し可能な媒体 1511 からインストールされる。このコンピュータプログラムが中央処理部（CPU）1501 により実行されると、本発明のシステムに限定される各種機能が実行される。

【0105】

なお、本発明の実施例に示すコンピュータ読み取り可能な媒体は、コンピュータ読み取り可能な信号媒体、コンピュータ読み取り可能な記憶媒体、又はこれらの何れかの組み合わせであってもよい。コンピュータ読み取り可能な記憶媒体は、例えば、電気、磁気、光、電磁、赤外線、又は半導体のシステム、機器又はデバイス、又は任意の以上の組み合わせであってもよいが、これらに限定されない。コンピュータ読み取り可能な記憶媒体のより具体的な例としては、1つ以上のリード線を有する電気接続、ポータブルコンピュータディスク、ハードディスク、ランダムアクセスメモリ（RAM）、読取り専用メモリ（ROM）、消去可能プログラマブル読取り専用メモリ（Erasable Programmable Read Only Memory：EPROM）、フラッシュメモリ、光ファイバ、ポータブルコンパクトディスク読取り専用メモリ（Compact Disc Read-Only Memory：CD-ROM）、光記憶装置、磁気記憶装置、又はこれらの任意の適切な組み合わせが挙げられるが、これらに限定されない。本発明では、コンピュータ読み取り可能な記憶媒体は、命令実行システム、機器又はデバイスによって使用されるか、又はそれらと組み合わせて使用されることができるプログラムを含むか、又は格納する任意の有形媒体であってもよい。本発明では、コンピュータ読み取り可能な信号媒体は、コンピュータ読み取り可能なプログラムコードが搬送されるベースバンド又は搬送波の一部として伝搬されるデータ信号を含んでもよい。そのような伝搬データ信号は、限定されるものではないが、電磁信号、光信号、又は上述の任意の適切な組み合わせを含む様々な形態であってもよい。コンピュータ読み取り可能な信号媒体は、コンピュータ読み取り可能な記憶媒体以外の任意のコンピュータ読み取り可能な媒体であってもよく、コンピュータ読み取り可能な媒体は、命令実行システム、機器又はデバイスによって使用するために、又はそれらと組み合わせて使用するために、プログラムを送信、伝播、又は送信してもよい。コンピュータ読み取り可能な媒体上に含まれるプログラムコードは、無線、有線等、又は上記の任意の適切な組み合わせを含むが、これらに限定されない任意の適切な媒体で送信してもよい。

【0106】

本発明の実施例の1つの態様では、コンピュータ読み取り可能な記憶媒体に記憶されたコンピュータ命令を含むコンピュータプログラム製品又はコンピュータプログラムであっても、コンピュータ装置のプロセッサは、コンピュータ読み取り可能な記憶媒体からコンピュータ命令を読み取り、該コンピュータ命令を実行することによって、前記コンピュータ装置に上記の各態様に記載の方法を実行させる、コンピュータプログラム製品又はコンピュータプログラムを提供する。

【0107】

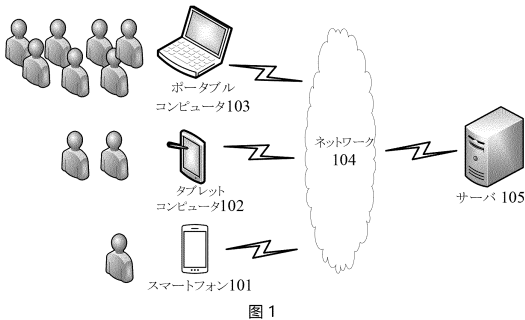
他の態様では、本発明は、コンピュータ読み取り可能な媒体をさらに提供する。該コンピュータ読み取り可能な媒体は、上記の実施例で説明した電子機器に含まれてもよいし、単独に存在してもよいし、該電子機器に組み込まれなくてもよい。上記のコンピュータ読み取り可能な媒体は、1つ又は複数のプログラムを保持し、上記の1つ又は複数のプログラムは、1つの電子機器により実行される際に、上記の実施例に記載の方法を電子機器に実現させる。

【0108】

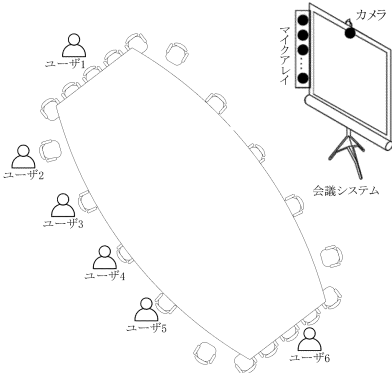
なお、本出願は、上述して図面に示した正確な構成に限定されるものではなく、その範囲から逸脱することなく種々の修正及び変更が可能である。本出願の範囲は、添付の特許請求の範囲によってのみ限定される。

【図面】

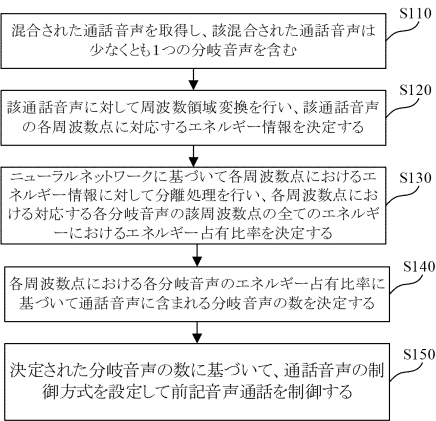
【図 1】



【図 2】



【図 3】



【図 4】



図 3

図 4

10

20

30

40

50

【図 5】

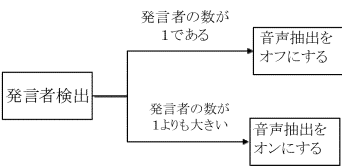


図 5

【図 6】

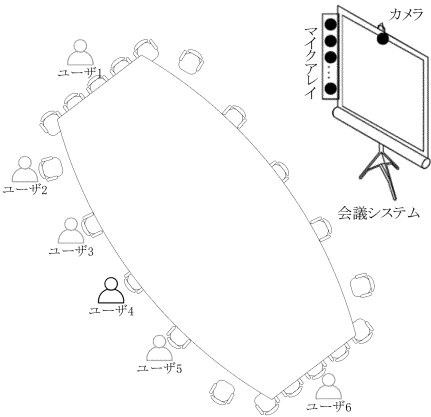


図 6

【図 7】

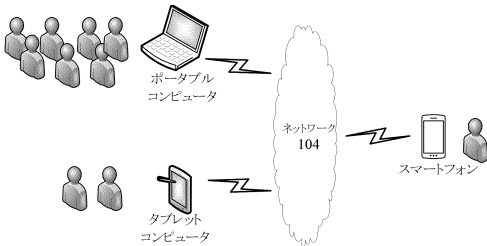


図 7

【図 8】

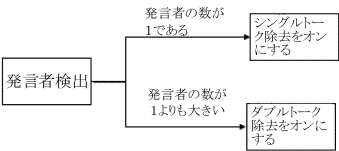


図 8

10

20

30

40

50

【図 9】

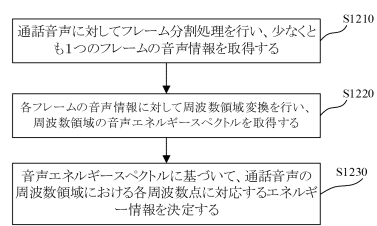


図 9

【図 10】

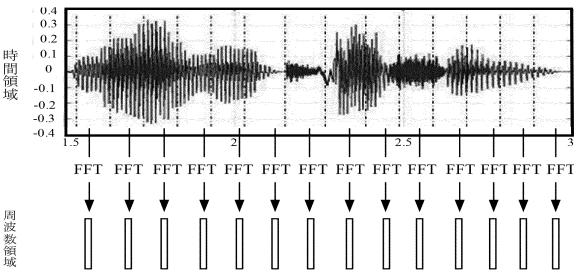


図 10

【図 11】

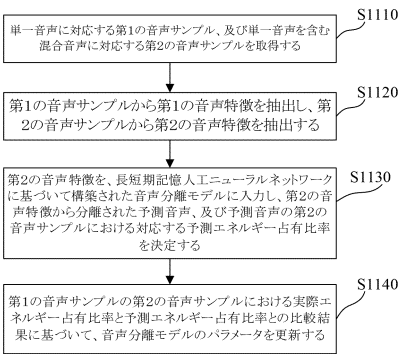


図 11

【図 12】

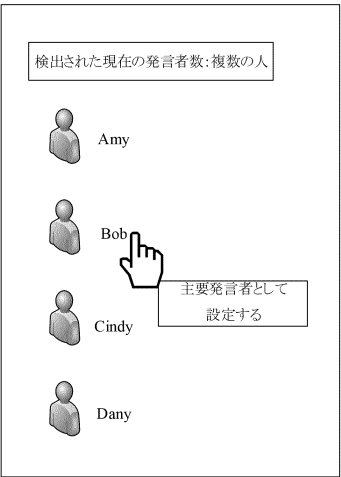


図 12

10

20

30

40

50

【図 1 3】

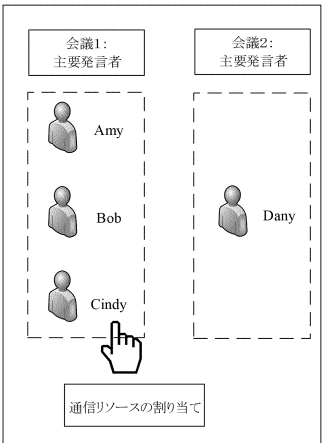


図 13

【図 1 4】

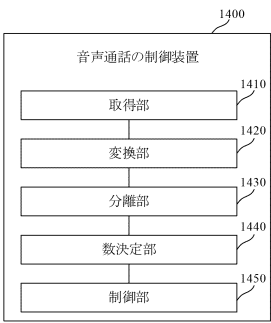


図 14

【図 1 5】

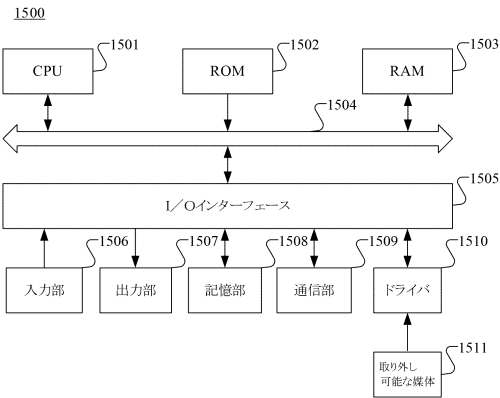


図 15

10

20

30

40

50

フロントページの続き

- (72)発明者

リー, ジュアンジュアン
中華人民共和国 518057 グアンドン シェンジェン ナンシャン・ディストリクト ミッドウ
エスト・ディストリクト・オブ・ハイテックパーク ケジジョンギ・ロード テンセント・ビルディ
ング 35エフ
- (72)発明者

シア, シャンジュン
中華人民共和国 518057 グアンドン シェンジェン ナンシャン・ディストリクト ミッドウ
エスト・ディストリクト・オブ・ハイテックパーク ケジジョンギ・ロード テンセント・ビルディ
ング 35エフ
- 審査官

中村 天真
- (56)参考文献

特開2018-205449 (JP, A)
特開2012-173584 (JP, A)
国際公開第2020/110228 (WO, A1)
特開2020-134657 (JP, A)
特開2010-066506 (JP, A)
- (58)調査した分野

(Int.Cl., DB名)
G10L 21/00-25/93
IEEE Xplore