

(19) 日本国特許庁 (JP)

(12) 特 許 公 報 (B2)

(11) 特許番号

特許第6911208号

(P6911208)

(45) 発行日 令和3年7月28日 (2021.7.28)

(24) 登録日 令和3年7月9日 (2021.7.9)

(51) Int. Cl. F I
G 1 O L 21/007 (2013.01) G 1 O L 21/007
G 1 O L 25/30 (2013.01) G 1 O L 25/30
G 1 O L 25/90 (2013.01) G 1 O L 25/90

請求項の数 27 (全 37 頁)

(21) 出願番号	特願2020-542648 (P2020-542648)	(73) 特許権者	507236292
(86) (22) 出願日	平成31年2月14日 (2019.2.14)		ドルビー ラボラトリーズ ライセンシン
(65) 公表番号	特表2021-508859 (P2021-508859A)		グ コーポレイション
(43) 公表日	令和3年3月11日 (2021.3.11)		アメリカ合衆国 94103 カリフォル
(86) 国際出願番号	PCT/US2019/017941		ニア州 サンフランシスコ マーケット
(87) 国際公開番号	W02019/161011		ストリート 1275
(87) 国際公開日	令和1年8月22日 (2019.8.22)	(74) 代理人	100107766
審査請求日	令和2年8月6日 (2020.8.6)		弁理士 伊東 忠重
(31) 優先権主張番号	18157080.5	(74) 代理人	100070150
(32) 優先日	平成30年2月16日 (2018.2.16)		弁理士 伊東 忠彦
(33) 優先権主張国・地域又は機関	欧州特許庁 (EP)	(74) 代理人	100135079
(31) 優先権主張番号	62/710,501		弁理士 宮崎 修
(32) 優先日	平成30年2月16日 (2018.2.16)		
(33) 優先権主張国・地域又は機関	米国 (US)		

最終頁に続く

(54) 【発明の名称】 発話スタイル転移

(57) 【特許請求の範囲】

【請求項1】

コンピュータで実装されるオーディオ処理方法であって：

発話合成器をトレーニングすることを含み、該トレーニングは：

(a) 一つまたは複数のプロセッサおよび一つまたは複数の非一時的記憶媒体を有する制御システムを介して実装される内容抽出プロセスによって、第一の人物の第一の発話に対応する第一のオーディオ・データを受領する段階と；

(b) 前記内容抽出プロセスによって、前記第一の発話に対応する第一のタイムスタンプ付けされた音素シーケンスおよび第一のピッチ輪郭データを生成する段階と；

(c) 前記制御システムを介して実装される第一のニューラルネットワークによって、前記第一のタイムスタンプ付けされた音素シーケンスおよび前記第一のピッチ輪郭データを受領する段階と；

(d) 前記第一のニューラルネットワークによって、前記第一のタイムスタンプ付けされた音素シーケンスおよび前記第一のピッチ輪郭データに対応する第一のニューラルネットワーク出力を生成する段階であって、前記第一のニューラルネットワーク出力は複数のフレーム・サイズを含む、段階と；

(e) 前記制御システムを介して実装される第二のニューラルネットワークによって、前記第一のニューラルネットワーク出力を受領する段階であって、前記第二のニューラルネットワークはモジュールの階層構成を含み、各モジュールは異なる時間分解能で動作し、前記第一のニューラルネットワークは、前記第一のニューラルネットワーク出力の前記

10

20

複数のフレーム・サイズのそれぞれが前記第二のニューラルネットワークのあるモジュールの時間分解能に対応するよう、前記第一のニューラルネットワーク出力を生成している、段階と；

(f) 前記第二のニューラルネットワークによって、第一の予測されたオーディオ信号を生成する段階と；

(g) 前記制御システムを介して、前記第一の予測されたオーディオ信号を第一の試験データと比較する段階であって、前記試験データは、前記第一の人物の発話に対応するオーディオ・データである、段階と；

(h) 前記制御システムを介して、前記第一の予測されたオーディオ信号についての損失関数値を決定する段階と；

(i) 前記第一の予測されたオーディオ信号についての現在の損失関数値と前記第一の予測されたオーディオ信号についての以前の損失関数値との間の差が所定の値以下になるまで(a)ないし(h)を繰り返す段階であって、(f)を繰り返すことは、前記第二のニューラルネットワークの少なくとも一つの重みに対応する少なくとも一つの非一時的記憶媒体位置の物理的状態を変更することを含む、段階とを含む、オーディオ処理方法。

【請求項2】

(a)は、前記第一の人物の前記第一の発話に対応する第一のタイムスタンプ付けされたテキストを受領することをさらに含む、請求項1記載のオーディオ処理方法。

【請求項3】

(a)は、前記第一の人物に対応する第一の識別データを受領することをさらに含む、請求項1または2記載のオーディオ処理方法。

【請求項4】

前記発話合成器を発話生成のために制御することをさらに含み、該発話生成は：

(j) 前記内容抽出プロセスによって、第二の人物の第二の発話に対応する第二のオーディオ・データ、前記第二の発話に対応する第二のタイムスタンプ付けされたテキストおよび前記第一の人物の発話に対応する第一の識別データを受領する段階と；

(k) 前記内容抽出プロセスによって、前記第二の発話に対応する第二のタイムスタンプ付けされた音素シーケンスおよび第二のピッチ輪郭データを生成する段階と；

(l) 前記第一のニューラルネットワークによって、(k)の前記第二のタイムスタンプ付けされた音素シーケンスおよび前記第二のピッチ輪郭データを受領する段階と；

(m) 前記第一のニューラルネットワークによって、(k)の前記第二のタイムスタンプ付けされた音素シーケンスおよび前記第二のピッチ輪郭データに対応する第一のニューラルネットワーク出力を生成する段階と；

(n) 前記第二のニューラルネットワークによって、(m)の前記第一のニューラルネットワーク出力および前記第一の識別データを受領する段階と；

(o) 前記第二のニューラルネットワークによって、(m)の前記第一のニューラルネットワーク出力および前記第一の識別データに対応する合成されたオーディオ・データを生成する段階とを含む、

請求項3記載のオーディオ処理方法。

【請求項5】

合成されたオーディオ・データは、前記第一の人物の発話特性に従って前記第二の人物によって発声される単語に対応する、請求項4記載のオーディオ処理方法。

【請求項6】

前記トレーニングは、第一の言語で前記第一のオーディオ・データを受領することに関わり、前記合成されたオーディオ・データは、第二の言語で前記第二の人物によって発声された単語に対応する、請求項5記載のオーディオ処理方法。

【請求項7】

一つまたは複数のトランスデューサに、前記合成されたオーディオ・データを再生させることをさらに含む、請求項4ないし6のうちいずれか一項記載のオーディオ処理方法。

【請求項 8】

前記トレーニングは：

第三のニューラルネットワークによって、前記第一のオーディオ・データを受領する段階と；

前記第一の人物の発話に対応する第一の発話特性を決定し、エンコードされたオーディオ・データを出力するよう前記第三のニューラルネットワークをトレーニングする段階とをさらに含む、

請求項 4 ないし 7 のうちいずれか一項記載のオーディオ処理方法。

【請求項 9】

前記トレーニングは、前記エンコードされたオーディオ・データが前記第一の人物の発話に対応するかどうかを判定するよう第四のニューラルネットワークをトレーニングする段階をさらに含む、請求項 8 記載のオーディオ処理方法。

【請求項 10】

前記発話生成は：

前記第三のニューラルネットワークによって、前記第二のオーディオ・データを受領する段階と；

前記第三のニューラルネットワークによって、前記第二のオーディオ・データに対応する第二のエンコードされたオーディオ・データを生成する段階と；

前記第四のニューラルネットワークによって、前記第二のエンコードされたオーディオ・データを受領する段階と；

前記第四のニューラルネットワークが、修正された第二のエンコードされたオーディオ・データが前記第一の人物の発話に対応すると判定するまで、対話的プロセスを介して、修正された第二のエンコードされたオーディオ・データを生成し、前記第四のニューラルネットワークが、修正された第二のエンコードされたオーディオ・データが前記第一の人物の発話に対応すると判定した後、該修正された第二のエンコードされたオーディオ・データを前記第二のニューラルネットワークに提供する段階とをさらに含む、請求項 9 記載のオーディオ処理方法。

【請求項 11】

(a) ないし (h) を繰り返すことは、前記第一のニューラルネットワークまたは前記第二のニューラルネットワークの少なくとも一方を、現在の損失関数値に基づく逆方向伝搬を介してトレーニングすることに関わる、請求項 1 ないし 10 のうちいずれか一項記載のオーディオ処理方法。

【請求項 12】

前記第一のニューラルネットワークは、双方向再帰型ニューラルネットワークを含む、請求項 1 ないし 11 のうちいずれか一項記載のオーディオ処理方法。

【請求項 13】

インターフェース・システムと；一つまたは複数のプロセッサおよび該一つまたは複数のプロセッサに動作上結合された一つまたは複数の非一時的記憶媒体を有する制御システムとを有する発話合成装置であって、前記制御システムは、発話合成器を実装するよう構成されており、前記発話合成器は、内容抽出器、第一のニューラルネットワークおよび第二のニューラルネットワークを含み、前記第一のニューラルネットワークは双方向再帰型ニューラルネットワークを含み、前記第二のニューラルネットワークは階層構成をなすモジュールを含み、各モジュールは異なる時間分解能で動作し、前記第一のニューラルネットワークおよび前記第二のニューラルネットワークは：

(a) 前記インターフェース・システムを介して前記内容抽出器によって、第一の人物の第一の発話に対応する第一のオーディオ・データを受領する段階と；

(b) 前記内容抽出器によって、前記第一の発話に対応する第一のタイムスタンプ付けされた音素シーケンスおよび第一のピッチ輪郭データを生成する段階と；

(c) 前記第一のニューラルネットワークによって、前記第一のタイムスタンプ付けされた音素シーケンスおよび前記第一のピッチ輪郭データを受領する段階と；

(d) 前記第一のニューラルネットワークによって、前記第一のタイムスタンプ付けされた音素シーケンスおよび前記第一のピッチ輪郭データに対応する第一のニューラルネットワーク出力を生成する段階であって、前記第一のニューラルネットワーク出力は複数のフレーム・サイズを含み、各フレーム・サイズは前記第二のニューラルネットワークのあるモジュールの時間分解能に対応する、段階と；

(e) 前記第二のニューラルネットワークによって、前記第一のニューラルネットワーク出力を受領する段階と；

(f) 前記第二のニューラルネットワークによって、第一の予測されたオーディオ信号を生成する段階と；

(g) 前記第一の予測されたオーディオ信号を第一の試験データと比較する段階であって、前記試験データは前記第一の人物の発話に対応するオーディオ・データである、段階と；

(h) 前記第一の予測されたオーディオ信号についての損失関数値を決定する段階と；

(i) 前記第一の予測されたオーディオ信号についての現在の損失関数値と前記第一の予測されたオーディオ信号についての以前の損失関数値との間の差が所定の値以下になるまで (a) ないし (h) を繰り返す段階とを含むプロセスに従ってトレーニングされており、

前記制御システムは、発話生成のために前記発話合成器モジュールを制御するよう構成されており、前記発話生成は：

(j) 前記内容抽出器によって、前記インターフェース・システムを介して、第二の人物の第二の発話に対応する第二のオーディオ・データ、前記第二の発話に対応する第二のタイムスタンプ付けされたテキストおよび前記第一の人物の発話に対応する第一の識別データを受領する段階と；

(k) 前記内容抽出器によって、前記第二の発話に対応する第二のタイムスタンプ付けされた音素シーケンスおよび第二のピッチ輪郭データを生成する段階と；

(l) 前記第一のニューラルネットワークによって、(k) の前記第二のタイムスタンプ付けされた音素シーケンスおよび前記第二のピッチ輪郭データを受領する段階と；

(m) 前記第一のニューラルネットワークによって、(k) の前記第二のタイムスタンプ付けされた音素シーケンスおよび前記第二のピッチ輪郭データに対応する第一のニューラルネットワーク出力を生成する段階と；

(n) 前記第二のニューラルネットワークによって、(m) の前記第一のニューラルネットワーク出力および前記第一の識別データを受領する段階と；

(o) 前記第二のニューラルネットワークによって、(m) の前記第一のニューラルネットワーク出力および前記第一の識別データに対応する合成されたオーディオ・データを生成する段階とを含む、
発話合成装置。

【請求項 14】

前記合成されたオーディオ・データは、前記第一の人物の発話特性に従って前記第二の人物によって発声される単語に対応する、請求項 13 記載の発話合成装置。

【請求項 15】

前記トレーニングは、第一の言語で前記第一のオーディオ・データを受領することに関わり、合成されたオーディオ・データは、第二の言語で前記第二の人物によって発声された単語に対応する、請求項 14 記載の発話合成装置。

【請求項 16】

前記制御システムは、一つまたは複数のトランスデューサに、第二の合成されたオーディオ・データを再生させるよう構成される、請求項 13 ないし 15 のうちいずれか一項記載の発話合成装置。

【請求項 17】

合成されたオーディオ・データを生成することは、前記第二のニューラルネットワークの少なくとも一つの重みに対応する少なくとも一つの非一時的記憶媒体位置の物理的状态

10

20

30

40

50

を変更することを含む、請求項 13 ないし 16 のうちいずれか一項記載の発話合成装置。

【請求項 18】

インターフェース・システムと；

一つまたは複数のプロセッサおよび該一つまたは複数のプロセッサに動作上結合された一つまたは複数の非一時的記憶媒体を有する制御システムとを有する発話合成装置であって、

前記制御システムは、発話合成器を実装するよう構成されており、前記発話合成器は、内容抽出器、第一のニューラルネットワークおよび第二のニューラルネットワークを含み、前記第一のニューラルネットワークは双方向再帰型ニューラルネットワークを含み、前記第二のニューラルネットワークは階層構成をなすモジュールを含み、各モジュールは異なる時間分解能で動作しし、前記第二のニューラルネットワークは、請求項 1 ないし 12 のうちいずれか一項記載のオーディオ処理方法によって第一の話者の第一の発話に対応する第一の合成されたオーディオ・データを生成するようトレーニングされており、前記制御システムは：

(a) 前記内容抽出器によって前記インターフェース・システムを介して、第二の人物の第二の発話に対応する第二のオーディオ・データ、前記第二の発話に対応する第二のタイムスタンプ付けされたテキストおよび前記第一の人物の発話に対応する第一の識別データを受領する段階と；

(b) 前記内容抽出器によって、前記第二の発話に対応する第二のタイムスタンプ付けされた音素シーケンスおよび第二のピッチ輪郭データを生成する段階と；

(c) 前記第一のニューラルネットワークによって、(b) の前記第二のタイムスタンプ付けされた音素シーケンスおよび前記第二のピッチ輪郭データを受領する段階と；

(d) 前記第一のニューラルネットワークによって、(b) の前記第二のタイムスタンプ付けされた音素シーケンスおよび前記第二のピッチ輪郭データに対応する第一のニューラルネットワーク出力を生成する段階であって、前記第一のニューラルネットワーク出力は複数のフレーム・サイズを含み、各フレーム・サイズは前記第二のニューラルネットワークのあるモジュールの時間分解能に対応する、段階と；

(e) 前記第二のニューラルネットワークによって、(d) の前記第一のニューラルネットワーク出力および前記第一の識別データを受領する段階と；

(f) 前記第二のニューラルネットワークによって、(d) の前記第一のニューラルネットワーク出力および前記第一の識別データに対応する第二の合成されたオーディオ・データを生成する段階とを実行するよう前記発話合成器を制御するよう構成されている、発話合成装置。

【請求項 19】

前記第二の合成されたオーディオ・データは、前記第一の人物の発話特性に従って前記第二の人物によって発声される単語に対応する、請求項 18 記載の発話合成装置。

【請求項 20】

前記トレーニングは、第一の言語で前記第一のオーディオ・データを受領することに関わり、前記第二の合成されたオーディオ・データは、第二の言語で前記第二の人物によって発声された単語に対応する、請求項 19 記載の発話合成装置。

【請求項 21】

一つまたは複数のトランスデューサに、前記第二の合成されたオーディオ・データを再生させることをさらに含む、請求項 18 ないし 20 のうちいずれか一項記載の発話合成装置。

【請求項 22】

前記合成されたオーディオ・データを生成することは、前記第二のニューラルネットワークの少なくとも一つの重みに対応する少なくとも一つの非一時的記憶媒体位置の物理的状态を変更することを含む、請求項 18 ないし 21 のうちいずれか一項記載の発話合成装置。

【請求項 23】

インターフェース・システムと；一つまたは複数のプロセッサおよび該一つまたは複数のプロセッサに動作上結合された一つまたは複数の非一時的記憶媒体を有する制御システムとを有する発話合成装置であって、前記制御システムは、発話合成器を実装するよう構成されており、前記発話合成器は、内容抽出器、第一のニューラルネットワークおよび第二のニューラルネットワークを含み、前記第一のニューラルネットワークは双方向再帰型ニューラルネットワークを含み、前記第二のニューラルネットワークは階層構成をなすモジュールを含み、各モジュールは異なる時間分解能で動作し、前記第一のニューラルネットワークおよび前記第二のニューラルネットワークは：

(a) 前記インターフェース・システムを介して前記内容抽出器によって、目標話者の第一の発話に対応する第一のオーディオ・データを受領する段階と；

10

(b) 前記内容抽出器によって、前記第一の発話に対応する第一のタイムスタンプ付けされた音素シーケンスおよび第一のピッチ輪郭データを生成する段階と；

(c) 前記第一のニューラルネットワークによって、前記第一のタイムスタンプ付けされた音素シーケンスおよび前記第一のピッチ輪郭データを受領する段階と；

(d) 前記第一のニューラルネットワークによって、前記第一のタイムスタンプ付けされた音素シーケンスおよび前記第一のピッチ輪郭データに対応する第一のニューラルネットワーク出力を生成する段階であって、前記第一のニューラルネットワーク出力は複数のフレーム・サイズを含み、各フレーム・サイズは前記第二のニューラルネットワークのあるモジュールの時間分解能に対応する、段階と；

(e) 前記第二のニューラルネットワークによって、前記第一のニューラルネットワーク出力を受領する段階と；

20

(f) 前記第二のニューラルネットワークによって、第一の予測されたオーディオ信号を生成する段階と；

(g) 前記第一の予測されたオーディオ信号を第一の試験データと比較する段階と；

(h) 前記第一の予測されたオーディオ信号についての損失関数値を決定する段階と；

(i) 前記第一の予測されたオーディオ信号についての現在の損失関数値と前記第一の予測されたオーディオ信号についての以前の損失関数値との間の差が所定の値以下になるまで (a) ないし (h) を繰り返す段階とを含むプロセスに従ってトレーニングされており、

前記制御システムは、発話生成のために前記発話合成器モジュールを制御するよう構成されており、前記発話生成は：

30

(j) 前記内容抽出器によって、前記インターフェース・システムを介して、源話者の第二の発話に対応する第二のオーディオ・データ、前記第二の発話に対応する第二のタイムスタンプ付けされたテキストおよび前記目標話者の発話に対応する第一の識別データを受領する段階と；

(k) 前記内容抽出器によって、前記第二の発話に対応する第二のタイムスタンプ付けされた音素シーケンスおよび第二のピッチ輪郭データを生成する段階と；

(l) 前記第一のニューラルネットワークによって、(k) の前記第二のタイムスタンプ付けされた音素シーケンスおよび前記第二のピッチ輪郭データを受領する段階と；

(m) 前記第一のニューラルネットワークによって、(k) の前記第二のタイムスタンプ付けされた音素シーケンスおよび前記第二のピッチ輪郭データに対応する第一のニューラルネットワーク出力を生成する段階と；

40

(n) 前記第二のニューラルネットワークによって、(m) の前記第一のニューラルネットワーク出力および前記第一の識別データを受領する段階と；

(o) 前記第二のニューラルネットワークによって、(m) の前記第一のニューラルネットワーク出力および前記第一の識別データに対応する合成されたオーディオ・データを生成する段階とを含む、
発話合成装置。

【請求項 24】

前記合成されたオーディオ・データが、前記目標話者の発話特性に従って前記源話者に

50

よって発声された単語に対応する、請求項 2 3 記載の発話合成装置。

【請求項 2 5】

前記目標話者および前記源話者が、異なる年齢における同じ人物である、請求項 2 3 記載の発話合成装置。

【請求項 2 6】

前記目標話者の前記第一の発話は、第一の年齢におけるまたは該第一の年齢を含む年齢範囲の間の、ある人物の発話に対応し、前記源話者の前記第二の発話は、第二の年齢におけるその人物の発話に対応する、請求項 2 3 記載の発話合成装置。

【請求項 2 7】

前記第一の年齢が前記第二の年齢より若い年齢である、請求項 2 6 記載の発話合成装置 10

【発明の詳細な説明】

【技術分野】

【0001】

関連出願への相互参照

本願は2018年2月16日に提出された米国仮特許出願第62/710,501号および2019年1月28日に提出された第62/797,864号ならびに2019年2月28日に提出された欧州特許出願第1815708 0.5号の利益を主張するものである。各出願の内容は参照によってその全体において組み込まれる。

【0002】

20

技術分野

本開示はオーディオ信号の処理に関する。特に、本開示は発話スタイル転移実装のためのオーディオ信号の処理に関する。

【背景技術】

【0003】

人物Aの発話を人物Bのスタイルで現実的に提示することは困難である。人物Aと人物Bが異なる言語を話す場合に、困難はさらに増す。たとえば、英語の映画を標準中国語で吹き替える声優を考える。キャラクターAの声を当てる声優は、あたかもキャラクターAが標準中国語で話しているかのように発話を発生する必要がある。この場合、発話スタイル転移の目標は、標準中国語での声優の声を入力として使って、あたかもキャラクターAが流暢 30
な標準中国語を話しているかのように聞こえる声を生成することであろう。非特許文献 1 は、声変換のための、深双方向長短期メモリ・ベースの再帰型ニューラルネットワーク（Deep Bidirectional Long Short-Term Memory based Recurrent Neural Networks (DBLSTM-RNNs)）の使用を記載している。

【非特許文献 1】Lifa Sun et al., “Voice Conversion Using Deep Bidirectional Long Short-Term Memory Based Recurrent Neural Networks”, 2015 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP), IEEE, 19 April 2015, pp.4869-4873

【発明の概要】

【発明が解決しようとする課題】

40

【0004】

本稿ではさまざまなオーディオ処理方法が開示される。いくつかのそのような方法は、発話合成器をトレーニングすることに関わってもよい。いくつかの例では、方法がコンピュータ実装されてもよい。たとえば、方法は、少なくとも部分的には、一つまたは複数のプロセッサおよび一つまたは複数の非一時的記憶媒体を有する制御システムを介して実装されてもよい。いくつかのそのような例では、トレーニングは：（a）一つまたは複数のプロセッサおよび一つまたは複数の非一時的記憶媒体を有する制御システムを介して実装される内容抽出プロセスによって、第一の人物の第一の発話に対応する第一のオーディオ・データを受領する段階と；（b）前記内容抽出プロセスによって、前記第一の発話に対応する第一のタイムスタンプ付けされた音素シーケンスおよび第一のピッチ輪郭データを 50

生成する段階と；（c）前記制御システムを介して実装される第一のニューラルネットワークによって、前記第一のタイムスタンプ付けされた音素シーケンスおよび前記第一のピッチ輪郭データを受領する段階と；（d）前記第一のニューラルネットワークによって、前記第一のタイムスタンプ付けされた音素シーケンスおよび前記第一のピッチ輪郭データに対応する第一のニューラルネットワーク出力を生成する段階であって、前記第一のニューラルネットワーク出力は複数のフレーム・サイズを含む、段階と；（e）前記制御システムを介して実装される第二のニューラルネットワークによって、前記第一のニューラルネットワーク出力を受領する段階であって、前記第二のニューラルネットワークはモジュールの階層構成を含み、各モジュールは異なる時間分解能で動作する、段階と；（f）前記第二のニューラルネットワークによって、第一の予測されたオーディオ信号を生成する段階と；（g）前記制御システムを介して、前記第一の予測されたオーディオ信号を第一の試験データと比較する段階と；（h）前記制御システムを介して、前記第一の予測されたオーディオ信号についての損失関数値を決定する段階と；（i）前記第一の予測されたオーディオ信号についての現在の損失関数値と前記第一の予測されたオーディオ信号についての以前の損失関数値との間の差が所定の値以下になるまで（a）ないし（h）を繰り返す段階とに関わってもよい。

10

【0005】

いくつかの実装によれば、当該方法の少なくともいくつかの動作は、少なくとも一つの非一時的記憶媒体位置の物理的状态を変更することに関わってもよい。たとえば、（f）を繰り返すことは、前記第二のニューラルネットワークの少なくとも一つの重みと対応する少なくとも一つの非一時的記憶媒体位置の物理的状态を変更することに関わってもよい。

20

【0006】

いくつかの例では、（a）は、第一の人物の第一の発話に対応する第一のタイムスタンプ付けされたテキストを受領することに関わってもよい。代替的または追加的に、（a）は、第一の人物に対応する第一の識別データを受領することに関わってもよい。

【0007】

いくつかの実装では、本方法は、発話合成器を発話生成のために制御することに関わってもよい。いくつかのそのような実施形態では、発話生成は：（j）前記内容抽出プロセスによって、第二の人物の第二の発話に対応する第二のオーディオ・データ、前記第二の発話に対応する第二のタイムスタンプ付けされたテキストおよび第一の人物の発話に対応する第一の識別データを受領する段階と；（k）前記内容抽出プロセスによって、前記第二の発話に対応する第二のタイムスタンプ付けされた音素シーケンスおよび第二のピッチ輪郭データを生成する段階と；（l）前記第一のニューラルネットワークによって、（k）の前記第二のタイムスタンプ付けされた音素シーケンスおよび前記第二のピッチ輪郭データを受領する段階と；（m）前記第一のニューラルネットワークによって、（k）の前記第二のタイムスタンプ付けされた音素シーケンスおよび前記第二のピッチ輪郭データに対応する第一のニューラルネットワーク出力を生成する段階と；（n）前記第二のニューラルネットワークによって、（m）の前記第一のニューラルネットワーク出力および前記第一の識別データを受領する段階と；（o）前記第二のニューラルネットワークによって、（m）の前記第一のニューラルネットワーク出力および前記第一の識別データに対応する合成されたオーディオ・データを生成する段階とに関わってもよい。いくつかの例によれば、前記第二のオーディオ・データは、第一の人物が異なる年齢であったときの第一の人物の発話に対応してもよい。いくつかのそのような例では、前記第二のオーディオ・データは、現時点または最近の時点において第一の人物から受領される発話であってもよい。第一の人物の前記第一の発話は、たとえば、第一の人物がもっと若かったときの第一の人物の発話に対応してもよい。

30

40

【0008】

いくつかの例によれば、合成されたオーディオ・データは、第一の人物の発話特性に従って第二の人物によって発音される単語に対応してもよい。いくつかのそのような例では

50

、トレーニングは、第一の言語で前記第一のオーディオ・データを受領することに関わってもよく、合成されたオーディオ・データは、第二の言語で第二の人物によって発音された単語に対応してもよい。しかしながら、いくつかの代替例では、合成されたオーディオ・データは、第一の年齢の第一の人物の発話特性に従って第二の年齢で、または第一の人物が第一の年齢範囲内であった時の間に、第一の人物によって発音された単語に対応してもよい。第一の年齢は、たとえば、第二の年齢より若い年齢であってもよい。いくつかの例は、一つまたは複数のトランスデューサに、合成されたオーディオ・データを再生させることに関わってもよい。

【0009】

いくつかの実装によれば、前記トレーニングは、第三のニューラルネットワークによって、前記第一のオーディオ・データを受領することに関わってもよい。いくつかのそのような実装では、トレーニングは、第一の人物の発話に対応する第一の発話特性を決定し、エンコードされたオーディオ・データを出力するよう第三のニューラルネットワークをトレーニングすることに関わってもよい。

【0010】

いくつかの例では、前記トレーニングは、エンコードされたオーディオ・データが第一の人物の発話に対応するかどうかを判定するよう第四のニューラルネットワークをトレーニングすることに関わってもよい。いくつかのそのような例では、発話生成は：第三のニューラルネットワークによって、前記第二のオーディオ・データを受領する段階と；第三のニューラルネットワークによって、前記第二のオーディオ・データに対応する第二のエンコードされたオーディオ・データを生成する段階と；第四のニューラルネットワークによって、前記第二のエンコードされたオーディオ・データを受領する段階と；第四のニューラルネットワークが修正された第二のエンコードされたオーディオ・データが第一の人物の発話に対応すると判定するまで、対話的プロセスを介して、修正された第二のエンコードされたオーディオ・データを生成する段階とに関わってもよい。いくつかの実装によれば、第四のニューラルネットワークが、修正された第二のエンコードされたオーディオ・データが第一の人物の発話に対応すると判定した後、本方法は、修正された第二のエンコードされたオーディオ・データを第二のニューラルネットワークに提供することに関わってもよい。

【0011】

いくつかの実装では、（a）ないし（h）を繰り返すことは、第一のニューラルネットワークまたは第二のニューラルネットワークの少なくとも一方を、現在の損失関数値に基づく後方伝搬を介してトレーニングすることに関わってもよい。いくつかの例によれば、第一のニューラルネットワークは、双方向再帰型ニューラルネットワークを含んでいてもよい。

【0012】

本稿に記載される方法の一部または全部は、一つまたは複数の非一時的媒体上に記憶される命令（たとえばソフトウェア）に従って一つまたは複数の装置によって実行されてもよい。そのような非一時的な媒体は、ランダムアクセスメモリ（RAM）デバイス、読み出し専用メモリ（ROM）デバイスなどを含むがそれに限られない、本稿に記載されるもののようなメモリ・デバイスを含んでいてもよい。よって、本開示に記載される主題のさまざまな革新的な側面は、ソフトウェアが記憶されている非一時的媒体において実装されることができる。ソフトウェアは、たとえば、オーディオ・データを処理するよう少なくとも一つの装置を制御するための命令を含んでいてもよい。ソフトウェアは、たとえば、本稿に開示されるもののような制御システムの一つまたは複数のコンポーネントによって実行可能であってもよい。ソフトウェアは、たとえば、本稿に開示される方法のうちの一つまたは複数を実行するための命令を含んでいてもよい。

【0013】

本開示の少なくともいくつかの側面は、装置により実装されてもよい。たとえば、一つまたは複数のデバイスが、本稿に開示される方法を少なくとも部分的に実行するよう構成

10

20

30

40

50

されてもよい。いくつかの実装では、装置は、インターフェース・システムおよび制御システムを含んでいてもよい。インターフェース・システムは、一つまたは複数のネットワーク・インターフェース、前記制御システムとメモリ・システムとの間の一つまたは複数のインターフェース、前記制御システムと別のデバイスとの間の一つまたは複数のインターフェースおよび／または一つまたは複数の外部デバイス・インターフェースを含んでいてもよい。制御システムは、汎用単一チップもしくは複数のチップ・プロセッサ、デジタル信号プロセッサ（DSP）、特定用途向け集積回路（ASIC）、フィールドプログラマブルゲートアレイ（FPGA）または他のプログラム可能な論理デバイス、離散的なゲートもしくはトランジスタ論理または離散的なハードウェア・コンポーネントのうちの少なくとも一つを含みうる。よって、いくつかの実装では、制御システムは、一つまたは複数のプロセッサと、前記一つまたは複数のプロセッサに動作上結合された一つまたは複数の非一時的記憶媒体を含んでいてもよい。

10

【0014】

いくつかのそのような例によれば、本装置は、インターフェース・システムおよび制御システムを含んでいてもよい。制御システムは、たとえば、発話合成器を実装するよう構成されてもよい。いくつかの実装では、発話合成器は、内容抽出器、第一のニューラルネットワークおよび第二のニューラルネットワークを含んでいてもよい。第一のニューラルネットワークはたとえば、双方向再帰型ニューラルネットワークを含んでいてもよい。いくつかの実装によれば、第二のニューラルネットワークは、複数のモジュールを含んでいてもよく、該モジュールはいくつかの事例では階層構成をなすモジュールであってもよい。いくつかのそのような例では、各モジュールは、異なる時間分解能で動作してもよい。

20

【0015】

いくつかのそのような例によれば、第一のニューラルネットワークおよび第二のニューラルネットワークは：（a）前記インターフェース・システムを介して前記内容抽出器によって、第一の人物の第一の発話に対応する第一のオーディオ・データを受領する段階と；（b）前記内容抽出器によって、前記第一の発話に対応する第一のタイムスタンプ付けされた音素シーケンスおよび第一のピッチ輪郭データを生成する段階と；（c）前記第一のニューラルネットワークによって、前記第一のタイムスタンプ付けされた音素シーケンスおよび前記第一のピッチ輪郭データを受領する段階と；（d）前記第一のニューラルネットワークによって、前記第一のタイムスタンプ付けされた音素シーケンスおよび前記第一のピッチ輪郭データに対応する第一のニューラルネットワーク出力を生成する段階であって、前記第一のニューラルネットワーク出力は複数のフレーム・サイズを含み、各フレーム・サイズは前記第二のニューラルネットワークのモジュールのある時間分解能に対応する、段階と；（e）前記第二のニューラルネットワークによって、前記第一のニューラルネットワーク出力を受領する段階と；（f）前記第二のニューラルネットワークによって、第一の予測されたオーディオ信号を生成する段階と；（g）前記第一の予測されたオーディオ信号を第一の試験データと比較する段階と；（h）前記第一の予測されたオーディオ信号についての損失関数値を決定する段階と；（i）前記第一の予測されたオーディオ信号についての現在の損失関数値と前記第一の予測されたオーディオ信号についての以前の損失関数値との間の差が所定の値以下になるまで（a）ないし（h）を繰り返す段階とを含むプロセスに従ってトレーニングされていてもよい。

30

40

【0016】

いくつかの実装では、前記制御システムは、発話生成のために前記発話合成器モジュールを制御するよう構成されてもよい。前記発話生成は、たとえば：（j）前記内容抽出器によって、前記インターフェース・システムを介して、第二の人物の第二の発話に対応する第二のオーディオ・データ、前記第二の発話に対応する第二のタイムスタンプ付けされたテキストおよび第一の人物の発話に対応する第一の識別データを受領する段階と；（k）前記内容抽出器によって、前記第二の発話に対応する第二のタイムスタンプ付けされた音素シーケンスおよび第二のピッチ輪郭データを生成する段階と；（l）前記第一のニューラルネットワークによって、（k）の前記第二のタイムスタンプ付けされた音素シーケ

50

ンスおよび前記第二のピッチ輪郭データを受領する段階と；（m）前記第一のニューラルネットワークによって、（k）の前記第二のタイムスタンプ付けされた音素シーケンスおよび前記第二のピッチ輪郭データに対応する第一のニューラルネットワーク出力を生成する段階と；（n）前記第二のニューラルネットワークによって、（m）の前記第一のニューラルネットワーク出力および前記第一の識別データを受領する段階と；（o）前記第二のニューラルネットワークによって、（m）の前記第一のニューラルネットワーク出力および前記第一の識別データに対応する合成されたオーディオ・データを生成する段階とに関わってもよい。いくつかの例によれば、前記第二のオーディオ・データは、第一の人物が異なる年齢または異なる年齢範囲内であったときの第一の人物の発話に対応してもよい。いくつかのそのような例では、前記第二のオーディオ・データは、現時点または最近の時点において第一の人物から受領される発話であってもよい。第一の人物の前記第一の発話は、たとえば、第一の人物がもっと若かったときの第一の人物の発話に対応してもよい。

10

【0017】

いくつかの例によれば、合成されたオーディオ・データは、第一の人物の発話特性に従って第二の人物によって発音される単語に対応してもよい。いくつかのそのような実装では、トレーニングは、第一の言語で前記第一のオーディオ・データを受領することに関わってもよく、合成されたオーディオ・データは、第二の言語で第二の人物によって発音された単語に対応してもよい。しかしながら、いくつかの代替例では、合成されたオーディオ・データは、第一の年齢での第一の人物の発話特性に従って第二の年齢で、または第一の年齢を含む年齢範囲（たとえば21ないし25歳、26ないし30歳など）の間に、第一の人物によって発音された単語に対応してもよい。第一の年齢は、たとえば、第二の年齢より若い年齢であってもよい。いくつかの例は、制御システムは、一つまたは複数のトランスデューサに、第二の合成されたオーディオ・データを再生させるよう構成されてもよい。いくつかの実装によれば、合成されたオーディオ・データを生成することは、前記第二のニューラルネットワークの少なくとも一つの重みと対応する少なくとも一つの非一時的記憶媒体位置の物理的状態を変更することに関わってもよい。

20

【0018】

本開示の少なくともいくつかの代替的な側面は、装置により実装されてもよい。いくつかのそのような例によれば、装置は、インターフェース・システムおよび制御システムを含んでいてもよい。制御システムは、たとえば、発話合成器を実装するよう構成されてもよい。いくつかの実装では、発話合成器は、内容抽出器、第一のニューラルネットワークおよび第二のニューラルネットワークを含んでいてもよい。第一のニューラルネットワークはたとえば、双方向再帰型ニューラルネットワークを含んでいてもよい。いくつかの実装によれば、第二のニューラルネットワークは、複数のモジュールを含んでいてもよく、該モジュールはいくつかの事例では階層構成をなすモジュールであってもよい。いくつかのそのような例では、各モジュールは、異なる時間分解能で動作してもよい。

30

【0019】

いくつかのそのような例によれば、第二のニューラルネットワークは、第一の話者の第一の発話に対応する第一の合成されたオーディオ・データを生成するようトレーニングされてもよい。いくつかの実装では、制御システムは：（a）前記内容抽出器によって前記インターフェース・システムを介して、第二の人物の第二の発話に対応する第二のオーディオ・データ、前記第二の発話に対応する第二のタイムスタンプ付けされたテキストおよび第一の人物の発話に対応する第一の識別データを受領する段階と；（b）前記内容抽出器によって、第二の発話に対応する第二のタイムスタンプ付けされた音素シーケンスおよび第二のピッチ輪郭データを生成する段階と；（c）第一のニューラルネットワークによって、（b）の前記第二のタイムスタンプ付けされた音素シーケンスおよび前記第二のピッチ輪郭データを受領する段階と；（d）前記第一のニューラルネットワークによって、（b）の前記第二のタイムスタンプ付けされた音素シーケンスおよび前記第二のピッチ輪郭データに対応する第一のニューラルネットワーク出力を生成する段階であって、前記第

40

50

一のニューラルネットワーク出力は複数のフレーム・サイズを含み、各フレーム・サイズは前記第二のニューラルネットワークのあるモジュールの時間分解能に対応する、段階と；（e）前記第二のニューラルネットワークによって、（d）の前記第一のニューラルネットワーク出力および前記第一の識別データを受領する段階と；（f）前記第二のニューラルネットワークによって、（d）の前記第一のニューラルネットワーク出力および前記第一の識別データに対応する第二の合成されたオーディオ・データを生成する段階とを実行するよう前記発話合成器を制御するよう構成されてもよい。

【0020】

いくつかの実装では、前記第二の合成されたオーディオ・データは、第一の人物の発話特性に従って第二の人物によって発音される単語に対応してもよい。いくつかのそのような例では、トレーニングは、第一の言語で前記第一のオーディオ・データを受領することに関わってもよく、合成されたオーディオ・データは、第二の言語で第二の人物によって発音された単語に対応してもよい。いくつかのそのような例は、一つまたは複数のトランスデューサに、合成されたオーディオ・データを再生させることに関わってもよい。いくつかの例では、合成されたオーディオ・データを生成することは、前記第二のニューラルネットワークの少なくとも一つの重みと対応する少なくとも一つの非一時的記憶媒体位置の物理的状态を変更することに関わってもよい。

【0021】

本明細書に記載される主題の一つまたは複数の実装の詳細は、付属の図面および下記の記述において記載される。他の特徴、側面および利点は、該記述、図面および請求項から明白になるであろう。下記の図面の相対的な寸法は同縮尺で描かれていないことがあることに注意されたい。さまざまな図面における同様の参照符号および記号は、一般に、同様の要素を示す。

【図面の簡単な説明】

【0022】

【図1】本稿に開示されるいくつかの実装による、発話スタイル転移のための一つまたは複数のニューラルネットワークをトレーニングするプロセスを示す。

【0023】

【図2】本稿に開示されるいくつかの実装による、発話スタイル転移のための一つまたは複数のトレーニングされたニューラルネットワークを使用するプロセスを示す。

【0024】

【図3】本稿に開示される方法の少なくともいくつかを実行するよう構成されうる装置のコンポーネントの例を示すブロック図である。

【0025】

【図4】一例による、発話合成器をトレーニングする方法の諸ブロックを概説する流れ図である。

【0026】

【図5】いくつかの例による、発話合成器トレーニング・システムの諸ブロックを示す。

【0027】

【図6】トレーニングされた発話合成器に合成されたオーディオ・データを生成させる例を示す。

【0028】

【図7】双方向再帰型ニューラルネットワークの一例を示す。

【0029】

【図8】声モデリング・ニューラルネットワークの例示的な諸ブロックおよび該声モデリング・ニューラルネットワークにトレーニングの間に提供されうる入力 of 例を示す。

【0030】

【図9】声モデリング・ニューラルネットワークの例示的な諸ブロックおよび該声モデリング・ニューラルネットワークに発話生成プロセスの間に提供されうる入力 of 例を示す。

【0031】

10

20

30

40

50

【図10】オートエンコーダの諸ブロックの例を示す。

【0032】

【図11】オートエンコーダを含む発話合成器をトレーニングするプロセスのための例示的な諸ブロックを示す。

【0033】

【図12A】話者分類器をトレーニングするプロセスの間に使用されうる諸ブロックの例を示す。

【0034】

【図12B】話者分類器の一例を示す。

【0035】

【図13】発話合成のために話者分類器およびオートエンコーダを使用することの例を与える。

【0036】

【図14】一例による、整形ニューラルネットワークおよび声モデリング・ニューラルネットワークの諸ブロックを示す。

【発明を実施するための形態】

【0037】

下記の記述は、本開示のいくつかの革新的な側面およびこれらの革新的な側面が実装されうるコンテキストの例を記述する目的のためのある種の実装に向けられる。しかしながら、本稿の教示は、さまざまな異なる仕方で適用できる。さらに、記載される実施形態は、多様なハードウェア、ソフトウェア、ファームウェアなどで実装されうる。たとえば、本願の諸側面は、少なくとも部分的には、装置、二つ以上のデバイスを含むシステム、方法、コンピュータ・プログラム・プロダクトなどで具現されうる。よって、本願の諸側面は、ハードウェア実施形態、ソフトウェア実施形態（ファームウェア、常駐ソフトウェア、マイクロコードなどを含む）および／またはソフトウェアおよびハードウェアの両側面を組み合わせる実施形態の形を取りうる。そのような実施形態は、本稿において「回路」、「モジュール」または「エンジン」と称されうる。本願のいくつかの側面は、コンピュータ可読プログラムコードが具現されている一つまたは複数の非一時的媒体において具現されるコンピュータ・プログラム・プロダクトの形を取りうる。そのような非一時的媒体は、たとえば、ハードディスク、ランダムアクセスメモリ（RAM）、読み出し専用メモリ（ROM）、消去可能なプログラム可能な読み出し専用メモリ（EPROMまたはフラッシュメモリ）、ポータブルなコンパクトディスク読み出し専用メモリ（CD-ROM）、光記憶デバイス、磁気記憶デバイスまたは上記の任意の好適な組み合わせを含みうる。よって、本開示の教示は、図面に示されるおよび／または本稿に記載される実装に限定されることは意図されておらず、広い適用可能性をもつ。

【0038】

発話スタイル転移（speech style transfer）は、時に、「声変成」または「声変換」と称され、よって、これらの用語は本稿では交換可能に使われることがある。図1は、本稿に開示されるいくつかの実装に従って、発話スタイル転移のために一つまたは複数のニューラルネットワークをトレーニングするプロセスを示している。図1に示される例では、トレーニング・プロセスは、内容抽出ブロックに入力される、ある人物（話者A）の発話に対応するオーディオ・データを提供することに関わる。話者Aは、いくつかの開示される例では、「目標話者（target speaker）」と称されてもよい。

【0039】

この例によれば、内容抽出ブロックは、テキスト入力を要求しない。ニューラルネットワーク・ベースの音素分類器が、たとえば、入力発話に対応するタイムスタンプ付けされた音素シーケンスを得るために、内容抽出ブロックにおいてトレーニングされ、使用されてもよい。この例では、話者Aの発話に対応するオーディオ・データだけが、入力として与えられるが、代替例では、人物Aの発話に対応するテキストが、入力発話波形と一緒に入力されてもよい。いくつかのそのような実装によれば、テキストは、タイムスタンプと

10

20

30

40

50

一緒に与えられてもよい。

【0040】

この例によれば、内容抽出ブロックの出力は、入力発話のピッチ輪郭 (pitch contour) (これは本稿では時に、「ピッチ輪郭データ」または「イントネーション」と称されることがある) と、入力発話に対応するタイムスタンプ付けされた音素シーケンス (これは時に発話テンポと称されることがある) とを含むデータである。ピッチ輪郭データは、たとえば、入力発話のピッチを時間を追って追跡する関数または曲線でありうる。いくつかの実装によれば、ある特定の時点におけるピッチ輪郭データは、その時点における入力発話の基本周波数、入力発話の基本周波数の対数值、入力発話の基本周波数の正規化された対数值またはその時点における入力発話のピーク (最高エネルギー) 周波数に対応してもよい。しかしながら、代替例では、ピッチ輪郭データは、所与の時点における複数のピッチに対応してもよい。

10

【0041】

この例によれば、ピッチ輪郭データおよびタイムスタンプ付けされた音素シーケンスは、入力発話について使用されている単語 (「トレーニング・データセット」) の同じ系列を生成するよう一つまたは複数のボカール・モデル・ニューラルネットワークをトレーニングするために使用される。好適なニューラルネットワークのいくつかの例は、後述する。いくつかの実装によれば、トレーニング・プロセスは、二つ以上のニューラルネットワークをトレーニングすることに関わってもよい。いくつかのそのような実装によれば、トレーニング・プロセスは、一つまたは複数のニューラルネットワークの出力を第二のニューラルネットワークに提供することに関わってもよい。第二のニューラルネットワークは、ボカール・モデル・ニューラルネットワークであってもよい。ニューラルネットワーク (単数または複数) がトレーニングされた後は、発話生成のために使用されうる。発話生成は、本稿では「発話合成」とも称される。

20

【0042】

図2は、本稿に開示されるいくつかの実装による発話スタイル転移のために一つまたは複数のトレーニングされたニューラルネットワークを使うプロセスを示している。この例では、別の人物 (話者B) の発話が、図1を参照して上述した同じ内容抽出ブロックに入力される。内容抽出ブロックによって出力される、第二の人物の発話の、ピッチ輪郭データおよびタイムスタンプ付けされた音素シーケンスは、話者Aの声のためにトレーニングされたボカール・モデル・ニューラルネットワークに提供される。いくつかの代替例では、ボカール・モデルニューラルネットワークは、第一の年齢の話者Aの声のためにトレーニングされている。別の人物の話者を内容抽出ブロックに提供する代わりに、いくつかの実装は、第二の年齢の話者Aの発話を提供することに関わる。第一の年齢は、たとえば、第二の年齢よりも若い年齢であってもよい。

30

【0043】

しかしながら、話者Aの声のためにトレーニングされたボカール・モデル・ニューラルネットワークは、話者Aに対応する、または話者Aの発話に対応する識別データ (これは単純な「ID」、またはより複雑な識別データでありうる) をも提供される。よって、この例によれば、ボカール・モデル・ニューラルネットワークは、話者Bの単語を、話者Bの発話テンポおよびイントネーションをもって、話者Aの声で出力する。別の言い方をすると、この例では、ボカール・モデル・ニューラルネットワーク出力の合成されたオーディオ・データは、ボカール・モデル・ニューラルネットワークによって学習された話者Aの発話特性に従って話者Bによって発音された単語を含む。いくつかの代替例では、第一の年齢の話者Aの声のためにトレーニングされたボカール・モデル・ニューラルネットワークが、第一の年齢の話者Aに対応する識別データを提供される。第二の年齢における話者Aの入力単語は、第一の年齢の話者Aの声で出力されうる。第一の年齢は、たとえば、第二の年齢よりも若い年齢であってもよい。

40

【0044】

いくつかの実装では、トレーニング・プロセスは、第一のオーディオ・データを第一の

50

言語で受領することに関わってもよく、合成されたオーディオ・データは、第二の言語で第二の人物によって発音される単語に対応してもよい。たとえば、話者Aが英語を話すことが知られており、話者Bが標準中国語を話すことが知られている場合、話者Aについての識別データは英語（言語1）と関連付けられてもよく、話者Bについての識別データは標準中国語（言語2）と関連付けられることができる。生成フェーズでは、ボークル・モデル・ニューラルネットワーク（単数または複数）は、人物Bからの発話を与えられるが、人物Aについての識別データを与えられる。結果は、人物Aのスタイルでの、言語2での発話である。

【0045】

よって、いくつかの実装は、言語2（たとえば英語）に存在しない言語1（たとえば標準中国語）における音について、ボークル・モデル・ニューラルネットワーク（単数または複数）をトレーニングすることに関わってもよい。いくつかのトレーニング実装は、英語について一つ、標準中国語について一つではなく、両言語についての合同音素集合（「上位集合」）を使用してもよい。いくつかの例は、通常は言語2で話す目標話者が、言語2に対応する音がない言語1の音素（たとえば標準中国語の音素）と対応する音を出すよう促されるトレーニング・プロセスに関わってもよい。目標話者は、たとえば、与えられる音をマイクロフォンに向かって繰り返してもよい。音は、目標話者の声を入力するために使われるマイクロフォンに拾われないよう、ヘッドフォンを介して与えられる。

【0046】

いくつかの代替例によれば、トレーニング・プロセスの間、目標話者は、自分の母語（言語1）の音素のみを含む音を与えてもよい。いくつかのそのような例では、発話生成プロセスは、言語2の発話を生成するとき、言語1からの最も類似した音素を使うことに関わってもよい。

【0047】

いくつかの代替的な実装では、トレーニング・プロセスは、原子的発声の上位集合の表現を生成することに関わってもよい。原子的発声は、人間が出すことのできる、音素および非発話発声の両方を含み、摩擦音および声門音を含む。原子的発声は、発声の基本単位と見なす者もある。

【0048】

満足のいく発話スタイル転移方法およびデバイスを開発しようとさまざまな以前の試みがなされてきた。たとえば、「SampleRNN」は、モンリオール学習プログラム協会（Montreal Institute for Learning Algorithms, MILA）によって開発されたエンドツーエンドのニューラル・オーディオ生成モデルであり、これは一時に一つのオーディオ・サンプルを生成する（非特許文献2参照）。このモデルは、長い時間期間にわたる時間シーケンスにおける変動の根底にある源を捕捉するために、自己回帰型（autoregressive）深層ニューラルネットワークおよびステートフルなニューラルネットワークを階層構造において組み合わせる。

【非特許文献2】SAMPLERNN: An Unconditional End-To-End Neural Audio Generation Model、会議論文として公開、International Conference on Learning Representations, Toulon, France, April 24-26, 2017

【0049】

SampleRNNは声変換分野における一歩前進であったが、SampleRNNはいくつかの欠点をもつ。たとえば、SampleRNNによって生成される発話信号は、もごもごした声であり、了解可能な発話ではない。さらに、俳優によって伝えられる感情は、自然に声にされることはできない：人は同じ感情を種々の仕方で声にできるので、テキストの感情を抽出する意味解析を使用しても、十分ではない。さらに、SampleRNNは、複数の目標話者を扱う機構を提供しない。

【0050】

いくつかの開示される実装は、次の潜在的な技術的改善のうちの一つまたは複数を提供する。たとえば、さまざまな開示される実装に従って生成される発話信号は、もごもごし

10

20

30

40

50

た声ではなく、了解可能な発話である。本稿に開示されるいくつかの実装によれば、複数の話者の声スタイルが、トレーニング・プロセスを通じて学習されてもよい。いくつかの実施形態は、源話者と目標話者からのパラレル発話をトレーニングのために要求しない。いくつかの実装は、生成される発話信号の韻律を改善するために、声優の入力発話信号（または他の語り手の入力発話信号）の内容的意味および／または入力発話信号の特性を考慮に入れる。いくつかの実装は、より了解可能な、自然に聞こえる発話信号を生成するために、複数の目標話者の学習可能な高次元表現を提供する。

【0051】

図3は、本稿に開示される方法の少なくともいくつかを実行するよう構成されうる装置のコンポーネントの例を示すブロック図である。いくつかの例では、装置305は、パーソナルコンピュータ、デスクトップコンピュータまたはオーディオ処理を提供するよう構成されている他のローカル装置であってもよく、またはそれを含んでいてもよい。いくつかの例では、装置305は、サーバーであってもよく、またはサーバーを含んでいてもよい。いくつかの例によれば、装置305は、ネットワーク・インターフェースを介したサーバーとの通信のために構成されたクライアント装置であってもよい。装置305のコンポーネントは、ハードウェアを介して、非一時的媒体に記憶されたソフトウェアを介して、ファームウェアを介しておよび／またはそれらの組み合わせによって、実装されてもよい。図3および本願で開示される他の図面に示されるコンポーネントの型および数は、単に例として示されている。代替的な実装は、より多数の、より少数のおよび／または異なるコンポーネントを含んでいてもよい。

【0052】

この例では、装置305は、インターフェース・システム310および制御システム315を含む。インターフェース・システム310は、一つまたは複数のネットワーク・インターフェース、制御システム315とメモリ・システムとの間の一つまたは複数のインターフェースおよび／または一つまたは複数の外部装置インターフェース（一つまたは複数のユニバーサルシリアルバス（USB）インターフェースのような）を含んでいてもよい。いくつかの実装では、インターフェース・システム310は、ユーザー・インターフェース・システムを含んでいてもよい。ユーザー・インターフェース・システムは、ユーザーから入力を受け取るよう構成されてもよい。いくつかの実装では、ユーザー・インターフェース・システムは、ユーザーにフィードバックを提供するよう構成されてもよい。たとえば、ユーザー・インターフェース・システムは、対応するタッチおよび／またはジェスチャー検出システムとともに一つまたは複数のディスプレイを含んでいてもよい。いくつかの例では、ユーザー・インターフェース・システムは、一つまたは複数のマイクロフォンおよび／またはスピーカーを含んでいてもよい。いくつかの例によれば、ユーザー・インターフェース・システムは、モーター、バイブレーターなどといった、触覚フィードバックを提供する装置を含んでいてもよい。制御システム315は、たとえば、汎用単一チップもしくは複数チップ・プロセッサ、デジタル信号プロセッサ（DSP）、特定用途向け集積回路（ASIC）、フィールドプログラマブルゲートアレイ（FPGA）または他のプログラム可能な論理デバイス、離散的なゲートもしくはトランジスタ論理および／または離散的なハードウェア・コンポーネントを含みうる。

【0053】

いくつかの例では、装置305は、単一の装置において実装されてもよい。しかしながら、いくつかの実装では、装置305は、二つ以上の装置において実装されてもよい。いくつかのそのような実装では、制御システム315の機能は、二つ以上の装置に含められてもよい。いくつかの例では、装置305は別の装置のコンポーネントであってもよい。

【0054】

図4は、一例による、発話合成器をトレーニングする方法の諸ブロックを概説する流れ図である。方法は、いくつかの事例では、図3の装置または別の型の装置によって実行されてもよい。いくつかの例では、方法400のブロックは、一つまたは複数の非一時的媒体に記憶されたソフトウェアを介して実装されてもよい。方法400のブロックは、本稿

に記載される他の方法と同様に、必ずしも示される順序で実行されない。さらに、そのような方法は、図示および／または記載されたものより多数または少数のブロックを含んでいてもよい。

【0055】

ここで、ブロック405は、第一の人物の発話に対応するデータを受領することに関わる。この例では、「第一の人物」は、目標話者であり、実際にその人の声のために発話合成器がトレーニングされている第一の人物であってもなくてもよい。一貫した用語を維持するために、ブロック405で受領されるオーディオ・データは本稿では「第一のオーディオ・データ」と称されてもよく、第一の人物の発話は本稿では「第一の発話」と称されてもよい。

10

【0056】

いくつかの例では、第一の人物からの第一のオーディオ・データのみが入力として提供されてもよい。そのような実装によれば、第一の発話からテキストが得られてもよい。しかしながら、そのような方法は、発話 - テキスト変換方法における潜在的な不正確さのために、最適ではないことがある。よって、代替の実装では、ブロック405は、第一の発話に対応するテキストを受領することに関わる。受領されるテキストは、第一のオーディオ・データの諸時刻に対応する諸タイムスタンプを含んでいてもよい。

【0057】

いくつかの例では、ブロック405は、第一の人物に対応する識別データを受領することに関わってもよい。該識別データは、本稿では「第一の識別データ」と称されてもよい。いくつかの例によれば、第一の識別データは、本質的には単に「これは話者Aである」ということを示しうる。しかしながら、代替例では、第一の識別データは、第一の発話の一つまたは複数の属性に関する情報を含んでいてもよい。下記にいくつかの例を記載する。トレーニングのコンテキストでは、識別データを提供することは、一つまたは複数のニューラルネットワークが、複数の目標話者からの発話を用いてトレーニングされることを許容できる。代替的または追加的に、トレーニングのコンテキストにおいて識別データを提供することは、その人が異なる年齢の時の同じ人からの発話を区別できる。単純な例では、第一の年齢（または第一の年齢を含む年齢範囲）の話者Aの発話はA1と指定されてもよく、第二の年齢（または第二の年齢を含む年齢範囲）の話者Aの発話はA2と指定されてもよい。

20

30

【0058】

この例によれば、ブロック405は、一つまたは複数のプロセッサおよび一つまたは複数の有体の記憶媒体を有する制御システムを介して実装される内容抽出プロセスによって、前記第一のオーディオ・データを受領することに関わる。内容抽出プロセスは、たとえば、一つまたは複数の非一時的記憶媒体に記憶されたソフトウェアに従って制御システムの一つまたは複数のプロセッサによって実装されてもよい。

【0059】

この例では、内容抽出プロセスは、ブロック410において、第一の発話に対応するタイムスタンプ付けされた音素シーケンスおよびピッチ輪郭データを生成するよう構成される。一貫した用語を維持するために、ブロック410で生成されるタイムスタンプ付けされた音素シーケンスは、本稿では、「第一のタイムスタンプ付けされた音素シーケンス」と称されてもよく、ブロック410で生成されたピッチ輪郭データは本稿では「第一のピッチ輪郭データ」と称されてもよい。

40

【0060】

この実装によれば、ブロック405は、内容抽出プロセスによって生成された第一のタイムスタンプ付けされた音素シーケンスおよび第一のピッチ輪郭データを、ニューラルネットワークによって、受領することに関わる。第一のニューラルネットワークは、本稿では「整形 (conditioning) ネットワーク」と称されてもよい。いくつかの実施形態では、第一のニューラルネットワークは、第一のタイムスタンプ付けされた音素シーケンスおよび第一のピッチ輪郭データを、声モデリング・ニューラルネットワークに提供される前に

50

前処理するまたは整えることがあるからである。

【0061】

ブロック420は、第一のニューラルネットワークがどのようにして第一のタイムスタンプ付けされた音素シーケンスおよび第一のピッチ輪郭データを整えるかの一例を与えている。この例では、第一のニューラルネットワークは、第一のタイムスタンプ付けされた音素シーケンスおよび第一のピッチ輪郭データに対応する第一のニューラルネットワーク出力を生成するよう構成される。ここで、該第一のニューラルネットワーク出力は、二つ以上（「複数」）のフレーム・サイズで生成される。第一のニューラルネットワーク出力の諸フレーム・サイズは、いくつかの例では、第二のニューラルネットワークの好適な諸入力フレーム・サイズに対応しうる。

10

【0062】

いくつかの例では、第一のニューラルネットワークは、第一のニューラルネットワーク出力を第二のニューラルネットワークに渡す前に、複数のフレームを処理してもよい。しかしながら、代替の実装では、第一のニューラルネットワークは、第一のニューラルネットワーク出力の複数のフレームが処理される間、第一のニューラルネットワーク出力を第二のニューラルネットワークに渡すのを遅らせなくてもよい。

【0063】

よって、この例では、ブロック425は、制御システムを介して実装される第二のニューラルネットワークによって、第一のニューラルネットワーク出力を受領することに関わる。この例では、第二のニューラルネットワークは、階層構造をなすモジュールを含み、各モジュールは異なる時間分解能で動作する。第一のニューラルネットワーク出力のフレーム・サイズは、第二のニューラルネットワークのモジュールの時間分解能に対応する。いくつかの詳細な例を後述する。

20

【0064】

この実装によれば、ブロック430は、第二のニューラルネットワークによって、第一の予測されたオーディオ信号を生成することに関わる。この例では、ブロック435は、制御システムを介して、第一の予測されたオーディオ信号を第一の試験データと比較することに関わり、ブロック440は、制御システムを介して、第一の予測されたオーディオ信号についての損失関数値を決定することに関わる。当業者によって理解されるであろうように、損失関数値は、第二のニューラルネットワークをトレーニングするために使用されうる。いくつかの実装によれば、損失関数値は、第一のニューラルネットワークをトレーニングするためにも使用されてもよい。

30

【0065】

いくつかの例によれば、第二のニューラルネットワークをトレーニングすること（およびいくつかの事例では第一のニューラルネットワークをトレーニングすること）は、損失関数が比較的「平坦」になり、現在の損失関数値と前の損失関数値（たとえば直前の損失関数値）との間の差が閾値以下になるまで続けられてもよい。図4に示される例では、ブロック445は、第一の予測されたオーディオ信号についての現在の損失関数値と第一の予測されたオーディオ信号についての前の損失関数値との間の差が所定の値、たとえば1.90、1.92、1.94、1.96、1.98、2.00など以下になるまで、ブロック405ないし440のうちの少なくともいくつかを反復することに関わる。後述するように、いくつかのブロックを反復すること（たとえばブロック420の反復および／またはブロック430の反復）は、第二のニューラルネットワークの少なくとも一つの重みと対応する少なくとも一つの有体な記憶媒体位置の物理的な状態を変化させることに関わってもよい。

40

【0066】

上記のように、識別データを提供することは、一つまたは複数のニューラルネットワークが複数の目標話者からの発話を用いてトレーニングされることを許容できる。図4を参照して上記したプロセスは、各目標話者について繰り返されてもよく、トレーニング結果が保存されて、各目標話者に対応する識別データと関連付けられてもよい。そのようなトレーニング結果は、たとえば、選択された目標話者に基づく発話生成のために使用されて

50

もよい。

【0067】

図5は、いくつかの例による発話合成器トレーニング・システムのブロックを示している。システム500は、いくつかの事例では、図3の装置または別の型の装置によって実行されてもよい。いくつかの例では、方法500のブロックは、一つまたは複数の非一時的媒体に記憶されたソフトウェアを介して実装されてもよい。代替的実装は、図示および／または記載されたものより多数または少数のブロックを含んでいてもよい。

【0068】

ここで、システム500は、目標話者の発話に対応する入力オーディオ・データを内容抽出ブロック505に提供するように構成される。目標話者は、いくつかの事例では、本稿の他所で言及される「第一の人物の第一の発話」に対応してもよい。いくつかの例では、入力オーディオ・データは、図3に示され上述したインターフェース・システム310のようなインターフェース・システムを介して提供されてもよい。いくつかの例では、入力オーディオ・データは、目標話者のその時の発話からのマイクロフォン信号を介して提供されてもよい。しかしながら、他の例では、入力オーディオ・データは、目標話者の前に記録された発話を介して提供されてもよい。いくつかのそのような例では、前に記録された発話は、入力オーディオ・データが内容抽出ブロック505に提供される時点より数年、あるいはさらには数十年前に記録されたものであってもよい。たとえば、前に記録された発話は、人の経歴のより早い段階の間に、たとえば俳優の経歴、音楽家の経歴、政治家の経歴などのより早い段階の特定の時間期間の間に記録されたものであってもよい。そのような以前に記録された発話はたとえば、映画またはテレビジョンのサウンドトラック、インタビュー録音などの一部に対応していてもよい。

【0069】

この例では、内容抽出ブロック505は、音素シーケンス整列推定器ブロック510およびピッチ輪郭推定器ブロック515を含む。この例では、音素シーケンス整列推定器ブロック510は、入力オーディオ・データおよび該入力オーディオ・データに対応するテキストを受領する。受領されたテキストは、いくつかの例では、入力オーディオ・データの諸時刻に対応するタイムスタンプを含んでいてもよい。いくつかの例によれば、受領されたテキストは、インタビュー文字起こしテキスト、スクリプト・テキスト、文字起こしテキストなどを含みうる。この例によれば、ピッチ輪郭推定器ブロック515は、入力オーディオ・データを受領するが、入力オーディオ・データに対応するテキストは受領しない。いくつかの例では、内容抽出ブロック505は、目標話者に対応する識別データをも受領してもよい。

【0070】

この例では、音素シーケンス整列推定器ブロック510は、入力オーディオ・データに対応するタイムスタンプ付けされた音素シーケンス512を生成するように構成される。いくつかの実装では、音素シーケンス整列推定器ブロック510は、入力オーディオ・データをメル周波数ケプストラル係数(mel-frequency cepstral coefficient、MFCC)に変換してもよい。MFCCは、音の短期パワースペクトルの表現であり、周波数の非線形なメル・スケール上での対数パワースペクトルの線形コサイン変換に基づく。音素シーケンス整列推定器ブロック510は、辞書を参照してテキストを既知の音素に変換するように構成されてもよい。音素シーケンス整列推定器ブロック510は、MFCC特徴と音素との間の整列を取るよう構成されてもよい。音素シーケンス整列推定器ブロック510は、いくつかの例では、Kaldi発話認識システムのような発話認識システムに基づく強制される整列器であってもよい。しかしながら、他の実装は、他の型の発話認識システムを用いてもよい。

【0071】

ここで、ピッチ輪郭推定器ブロック515は、入力オーディオ・データに対応するピッチ輪郭データ517を生成するように構成される。いくつかの例では、ピッチ輪郭データ517は、各オーディオ・フレームについて、ピッチの正規化された対数値を示してもよい。いくつかの例では、オーディオ・フレームは5ミリ秒(ms)の継続時間であってもよい。

。しかしながら、代替例は、4msのフレーム、6msのフレーム、8msのフレーム、10msのフレームなどといった、より小さな、またはより大きなオーディオ・フレームを実装してもよい。いくつかの例では、ピッチ輪郭推定器ブロック515は、ピッチ絶対値を表わす、たとえば200個の浮動小数点数をシーケンスで生成するよう構成されてもよい。ピッチ輪郭推定器ブロック515は、これらの浮動小数点数に対して対数演算を実行し、各目標話者について、結果として得られる値を正規化して、出力が絶対的なピッチ値（たとえば200.0Hz）ではなく、0.0のまわりの輪郭（たとえば0.5）となるようにするよう構成されてもよい。いくつかの例では、ピッチ輪郭推定器ブロック515は、発話時間特性を使用するよう構成されてもよい。いくつかの例によれば、ピッチ輪郭推定器ブロック515は、まず、異なるカットオフ周波数をもつ、いくつかの低域通過フィルタを使用してもよい。一例では、低域通過フィルタは、50Hzのカットオフ周波数を有していてもよく、よって、0から50Hzまでの間の信号を通過させる。他の実装は他のカットオフ周波数を有していてもよい。いくつかのそのような実装によれば、これらの低域通過フィルタのカットオフ周波数は、50Hzから500Hzの間であってもよい。フィルタ処理された信号が基本周波数のみからなる場合、ピッチ輪郭推定器ブロック515は、正弦波を形成しうる。基本周波数は、この正弦波の周期に基づいて得られてもよい。いくつかの事例では、ピッチ輪郭推定器ブロック515は、最良の基本周波数候補を選ぶために零交差およびピーク谷間隔を使用してもよい。いくつかの例では、ピッチ輪郭推定器ブロック515は、Githubで提供されているWorldピッチ推定器のようなピッチ推定器を含んでいてもよい。しかしながら、他の例では、ピッチ輪郭推定器ブロック515は、別の型のピッチ推定器を含んでいてもよい。

10

20

【0072】

いくつかの例では、内容抽出ブロック505は、現在のオーディオ・フレームが有声であるか無声であるかを示す指標を出力するよう構成される。いくつかの実装によれば、内容抽出ブロック505は、現在の音素、最も最近の諸音素のうちの一つまたは複数および一つまたは複数の将来の音素（たとえば次の音素および可能性としてはさらなる音素）を含む複数の音素を出力するよう構成される。あるそのような例によれば、内容抽出ブロック505は、現在の音素、二つの前の音素および次の二つの音素を含む五つの音素を出力するよう構成される。

【0073】

この実装によれば、タイムスタンプ付けされた音素シーケンス512およびピッチ輪郭データ517は、整形ニューラルネットワーク520によって受領される。整形ニューラルネットワーク520は、図4を参照して上述した「第一のニューラルネットワーク」の事例である。この例では、整形ニューラルネットワーク520は、タイムスタンプ付けされた音素シーケンス512およびピッチ輪郭データ517を、これらのデータが声モデリング・ニューラルネットワーク525に提供される前に、前処理するまたは整えるよう構成される。いくつかの例では、声モデリング・ニューラルネットワーク525は複数のモジュールを含んでいてもよく、各モジュールは異なるレートまたは時間分解能で動作する。整形ニューラルネットワーク520によって出力される諸フレーム・サイズは、声モデリング・ニューラルネットワーク525のモジュールの時間分解能に対応してもよい。

30

40

【0074】

この例では、声モデリング・ニューラルネットワーク525は、予測されたオーディオ信号530を生成し、該予測されたオーディオ信号530を損失関数決定ブロック535に提供するよう構成される。ここで、損失関数決定ブロック535は、予測されたオーディオ信号を試験データ540と比較し、予測されたオーディオ信号についての損失関数値を決定するよう構成される。この例によれば、試験データ540は目標話者の発話に対応するオーディオ・データである。いくつかの例では、試験データ540は、内容抽出ブロック505に以前に提供されていない、目標話者の発話に対応するオーディオ・データを含む。この例では、損失関数値は、声モデリング・ニューラルネットワーク525をトレーニングするために使用される。いくつかの実装によれば、損失関数値は、整形ニューラ

50

ルネットワーク520をトレーニングするためにも使用されてもよい。

【0075】

いくつかの例によれば、声モデリング・ニューラルネットワーク525をトレーニングすること（およびいくつかの事例では、整形ニューラルネットワーク520をトレーニングすること）は、損失関数が比較的「平坦」になり、現在の損失関数値と前の損失関数値との間の差が閾値以下になるまで続けられてもよい。のちに詳述するように、声モデリング・ニューラルネットワーク525をトレーニングすることは、声モデリング・ニューラルネットワーク525の重みおよび／または活性化関数と対応する有体な記憶媒体位置の物理的な状態を変化させることに関わってもよい。

【0076】

図6は、トレーニングされた発話合成器に、合成されたオーディオ・データを生成させる例を示している。この例では、声モデリング・ニューラルネットワーク525は、目標話者の声に対応するオーディオ・データを合成するようすでにトレーニングされている。この例によれば、源話者（source speaker）の発話に対応する入力オーディオ・データが内容抽出ブロック505に提供される。源話者の発話は、いくつかの事例では、本稿の他所で言及されている「第二の人物の第二の発話」に対応してもよい。しかしながら、他の例では、声モデリング・ニューラルネットワーク525は、すでに、話者が第一の年齢であったときまたは話者が第一の年齢を含む年齢範囲（たとえば21歳ないし25歳、26歳ないし30歳など）であった時間の間の目標話者の声に対応するオーディオ・データを合成するようトレーニングされていてもよい。いくつかのそのような例によれば、源話者の発話は、目標話者が異なる年齢であった時の間、たとえば目標話者がより高い年齢であったときの目標話者の発話に対応してもよい。

【0077】

この例では、音素シーケンス整列推定器ブロック510は、入力オーディオ・データおよび該入力オーディオ・データに対応するテキストを受領し、ピッチ輪郭推定器515は入力オーディオ・データを受領するが、入力オーディオ・データに対応するテキストは受領しない。受領されたテキストは、入力オーディオ・データの諸時刻に対応するタイムスタンプを含んでいてもよい。入力オーディオ・データが今は源話者に対応しているが、いくつかの実装では、内容抽出ブロック505は、目標話者に対応する識別データを受領する。源話者および目標話者が異なる年齢の同じ人物であるいくつかの実装では、所望される年齢または年齢範囲の目標話者についての識別データが、ブロック505に提供されてもよい。一つの単純な例では、第一の年齢（または年齢範囲）での話者Aの発話がA1と表わされ、第二の年齢（または年齢範囲）での話者Aの発話がA2と表わされてもよい。よって、そのような例では、システム500は、目標話者の発話特性に従って、源話者によって発声された単語に対応する合成されたオーディオ・データを生成する。

【0078】

この例では、音素シーケンス整列推定器ブロック510は、源話者からの入力オーディオ・データに対応するタイムスタンプ付けされた音素シーケンス512を生成するよう構成される。ここで、ピッチ輪郭推定器ブロック515は、源話者からの入力オーディオ・データに対応するピッチ輪郭データ517を生成するよう構成される。

【0079】

この実装によれば、整形ニューラルネットワーク520は、タイムスタンプ付けされた音素シーケンス512およびピッチ輪郭データ517を、これらのデータの整形されたバージョン（これが、本稿の他所で言及される「第一のニューラルネットワーク出力」の例である）が声モデリング・ニューラルネットワーク525に提供される前に前処理するまたは整えるように構成される。

【0080】

この例では、源話者からの入力オーディオ・データに対応するタイムスタンプ付けされた音素シーケンス512およびピッチ輪郭データ517の整形されたバージョンを受領することに加えて、声モデリング・ニューラルネットワーク525は、目標話者に対応する

10

20

30

40

50

識別データをも受領する。よって、この例では、声モデリング・ニューラルネットワーク 525 は、第一のニューラルネットワーク出力および第一の識別データに対応する合成されたオーディオ・データを含む予測されたオーディオ信号 530（「合成されたオーディオ・データ」ともいう）を生成するよう構成される。そのような事例では、合成されたオーディオ・データは、目標話者の発話特性に従って、源話者によって発声された単語に対応する。トレーニング・プロセスが第一の言語で目標話者に対応するオーディオ・データを受領することに関わるいくつかの例によれば、合成されたオーディオ・データは、第二の言語で源話者によって発声された単語に対応してもよい。

【0081】

この例では、声モデリング・ニューラルネットワーク 525 は、目標話者の発話に対応する発話を合成するようすでにトレーニングされているが、いくつかの実装では、予測されたオーディオ信号 530 は、記憶のため、一つまたは複数のトランスデューサを介した再生のためな出力される前に、評価され、洗練されてもよい。いくつかのそのような実装では、予測されるオーディオ信号 530 は損失関数決定ブロック 535 に提供されてもよい。損失関数決定ブロック 535 は、予測されたオーディオ信号 530 を試験データと比較してもよく、予測されたオーディオ信号 530 についての損失関数値を決定してもよい。損失関数値は、予測されたオーディオ信号 530 をさらに洗練するために使用されてもよい。いくつかの実装によれば、損失関数値は、整形ニューラルネットワーク 520 をトレーニングするために使用されてもよい。

【0082】

いくつかの実装によれば、第一のニューラルネットワークおよび／または第二のニューラルネットワークは再帰型ニューラルネットワークであってもよい。当業者に知られているように、再帰型ニューラルネットワークは、個々のユニットまたは「ニューロン」の間の接続が有向サイクルをなすニューラルネットワークの類型である。この特徴のため、再帰型ニューラルネットワークは、動的な時間的挙動を示す。フォードフォワード型ニューラルネットワークと異なり、再帰型ニューラルネットワークは、その内部メモリを、任意の入力シーケンスを処理するために使うことができる。この能力のため、再帰型ニューラルネットワークは、手書き認識または発話認識といったタスクに適用可能になる。

【0083】

基本的な再帰型ニューラルネットワークは、しばしば「ニューロン」と称されるネットワーク・ノードを含む。各ニューロンは、他のニューロンへの有向の（一方向の）接続をもつ。各ニューロンは、入力または入力の集合を与えられて、そのニューロンの出力を定義する、一般に「活性化」と称される、時間変化する、実数値の活性化関数をもつ。ニューロン間の各接続（「シナプス」とも称される）は、修正可能な実数値の重みをもつ。ニューロンは、入力ニューロン（ネットワーク外部からデータを受領する）、出力ニューロンまたは入力ニューロンから出力ニューロンへの途上でデータを修正する隠れニューロンでありうる。いくつかの再帰型ニューラルネットワークは、入力ニューロンの層と出力ニューロンの層の間に、隠れニューロンのいくつかの層を含んでいてもよい。

【0084】

ニューラルネットワークは、図 3 を参照して上記した制御システム 315 のような制御システムによって実装されてもよい。よって、第一のニューラルネットワークまたは第二のニューラルネットワークが再帰型ニューラルネットワークである実装については、第一のニューラルネットワークまたは第二のニューラルネットワークをトレーニングすることは、再帰型ニューラルネットワークにおける重みに対応する非一時的記憶媒体位置の物理的な状態を変更することに関わってもよい。前記記憶媒体位置は、制御システムによってアクセス可能なまたは制御システムの一部である一つまたは複数の記憶媒体の一部でありうる。上記の重みは、ニューロン間の接続に対応する。第一のニューラルネットワークまたは第二のニューラルネットワークをトレーニングすることも、ニューロンの活性化関数の値に対応する非一時的記憶媒体位置の物理的な状態を変更することに関わってもよい。

【0085】

第一のニューラルネットワークは、いくつかの例では、双方向再帰型ニューラルネットワークであってもよい。標準的な再帰型ニューラルネットワークでは、将来の時刻に対応する入力、現在の状態から到達できない。対照的に、双方向再帰型ニューラルネットワークは、その入力データが固定されることを要求しない。さらに、双方向再帰型ニューラルネットワークの将来の入力情報は、現在の状態から到達可能である。双方向再帰型ニューラルネットワークの基本的なプロセスは、反対の時間方向に対応する二つの隠れ層を同じ入力および出力に接続することである。この型の構造を実装することにより、双方向再帰型ニューラルネットワークの出力層におけるニューロンは、過去および将来の状態から情報を受領できる。双方向再帰型ニューラルネットワークは、入力のコンテキストが必要とされるときに特に有用である。たとえば、手書き認識アプリケーションでは、現在の文字の前後の文字の知識によって、パフォーマンスが向上される。

10

【0086】

図7は、双方向再帰型ニューラルネットワークの一例を示している。図7に示される層の数、各層におけるニューロンの数などは、単に例である。他の実装は、より多数またはより少数の層、各層内のニューロンなどを含んでいてもよい。

【0087】

図7では、ニューロン701は丸で表わされている。層705の「x」ニューロンは入力ニューロンであり、双方向再帰型ニューラルネットワーク700の外部からデータを受領するよう構成される。層730の「y」ニューロンは出力ニューロンであり、双方向再帰型ニューラルネットワーク700からデータを出力するよう構成される。層710~725におけるニューロンは、入力ニューロンから出力ニューロンへの途上でデータを修正する隠れニューロンである。いくつかの実装では、双方向再帰型ニューラルネットワーク700のニューロンは、シグモイド活性化関数、tanh活性化関数またはシグモイドおよびtanh活性化関数の両方を用いてもよい。四つの隠れ層が図7に示されているが、いくつかの実装は、より多数またはより少数の隠れ層を含んでいてもよい。いくつかの実装は、ずっと多くの隠れ層、たとえば数百または数千の隠れ層を含んでいてもよい。たとえば、いくつかの実装は、128、256、512、1024、2048またはより多くの隠れ層を含んでいてもよい。

20

【0088】

この例では、双方向再帰型ニューラルネットワーク700は、三列のニューロンを含み、各列は異なる時間に対応する。それらの時間はたとえば、入力データが双方向再帰型ニューラルネットワーク700に提供される時間区間に対応してもよい。中央の列704は、時間tに対応し、左の列702は時間t-1に対応し、右の列706は時間t+1に対応する。時間t-1は、たとえば、時間tの直前の時間に取りれたデータ・サンプルに対応してもよく、時間t+1は、たとえば、時間tの直後の時間に取りれたデータ・サンプルに対応してもよい。

30

【0089】

この例では、隠れ層710および715は反対の時間方向に対応する。隠れ層710のニューロン701はデータを時間的に前方に渡す。一方、隠れ層715のニューロンはデータを時間的に逆方向に渡す。しかしながら、隠れ層710は隠れ層715に入力を与えず、隠れ層715は隠れ層710に入力を与えない。

40

【0090】

特定の時間に対応する層705の入力ニューロン—たとえば時間tに対応する列704の入力ニューロン—は、隠れ層710のニューロンおよび隠れ層715のニューロンに情報を提供する。隠れ層710のニューロンおよび隠れ層715のニューロンは、同じ時間に対応する層720の単一のニューロンに情報を提供する。

【0091】

ニューラルネットワークの情報および処理の流れは入力ニューロンから出力ニューロンに進むが、図7は、点線矢印740によって描かれる、逆方向の「逆方向伝搬」（「逆伝搬」としても知られる）をも示している。逆伝搬は、データのバッチが処理された後に各

50

ニューロンの誤差寄与を計算するためにニューラルネットワークにおいて使用される方法である。逆伝搬は、損失関数の勾配を計算することによって、ニューロンの重みを調整するよう勾配降下最適化アルゴリズムを適用することに関わってもよい。

【0092】

図7に示した例では、層730の出力ニューロンは、損失関数決定ブロックに出力を提供してもよく、損失関数決定ブロックは層730の出力ニューロンに現在の損失関数値を提供してもよい。誤差が出力において計算され、点線矢印740によって示されるようにニューラルネットワークの諸層を通じて戻り方向に分配されるからである。

【0093】

この例では、逆伝搬は双方向再帰型ニューラルネットワークのコンテキストで図示され、記述されたが、逆伝搬技法は、他の型の再帰型ニューラルネットワークを含むがそれに限られない他の型のニューラルネットワークに適用されてもよい。たとえば、逆伝搬技法は、本稿の他所で記載される声モデリング・ニューラルネットワーク（「第二のニューラルネットワーク」）、オートエンコーダおよび／または発話分類器ニューラルネットワークに適用されてもよい。

【0094】

図8は、声モデリング・ニューラルネットワークの例示的なブロックおよびトレーニング中に声モデリング・ニューラルネットワークに与えられうる入力の例を示している。いくつかの例によれば、声モデリング・ニューラルネットワーク525のニューロンは、シグモイド活性化関数および／またはtanh活性化関数を用いてもよい。代替的または追加的に、声モデリング・ニューラルネットワーク525のニューロンは、整流線形ユニット（ReLU）活性化関数を用いてもよい。この例では、S_t、P_t、Q_tおよびF_tが声モデリング・ニューラルネットワーク525に提供される。ここで、S_tは目標話者識別データを表わし、P_tは目標話者について声モデリング・ニューラルネットワーク525によって生成された前の予測されたオーディオ信号を表わし、Q_tは、目標話者の声に対応する入力時間整列された音素シーケンスを表わし、F_tは、目標話者の声に対応する基本周波数輪郭データを表わす。

【0095】

この例では、声モデリング・ニューラルネットワーク525は、モジュール805、810、815を含み、そのそれぞれは異なる時間分解能で動作する。この例では、モジュール805は、モジュール810よりもフレーム当たり、より多くのサンプルを処理し、モジュール815はモジュール805よりもフレーム当たり、より多くのサンプルを処理する。いくつかのそのような例では、モジュール805は、フレーム当たりモジュール810の10倍のサンプルを処理し、モジュール815は、フレーム当たりモジュール810の80倍のサンプルを処理する。いくつかのそのような実装によれば、声モデリング・ニューラルネットワーク525は、上記のような、SampleRNNニューラルネットワークの修正されたバージョンを含んでいてもよい。SampleRNNニューラルネットワークはたとえば、複数の目標話者についてトレーニングされ、該複数の目標話者のうちの選択された一人に対応する合成されたオーディオ・データを生成するよう修正されてもよい。しかしながら、これらは単に例である。他の実装では、声モデリング・ニューラルネットワーク525は、異なる数のモジュールを含んでいてもよく、および／またはモジュールは異なるフレーム・サイズを処理するよう構成されてもよい。

【0096】

よって、異なるフレーム・サイズの入力データは、たとえば整形ニューラルネットワーク520（図8には示さず）によって声モデリング・ニューラルネットワーク525に入力されてもよい。一つのそのような例では、オーディオ・データは16kHzでサンプリングされ、よって、整形ニューラルネットワーク520は、モジュール805の各フレームについて、5msのオーディオ・データに相当する80サンプル（「大フレーム」サイズ）を提供してもよい。一つのそのような実装では、整形ニューラルネットワーク520は、モジュール810の各フレームについて、0.5msのオーディオ・データに相当する8サンプル（

10

20

30

40

50

「小フレーム」サイズ)を提供してもよい。一つのそのような実装では、整形ニューラルネットワーク520は、モジュール815の各フレームについて、40msのオーディオ・データに相当する640サンプル(「サンプル予測器」フレーム・サイズ)を提供してもよい。いくつかのそのような例では、モジュール810は、モジュール805の10倍の速さ、モジュール815の80倍の速さで動作してもよい。

【0097】

いくつかの実装では、整形ニューラルネットワーク520は、モジュール810に提供される同じ8個のサンプルを10回繰り返して、モジュール805への入力のための80サンプルを生成してもよい。いくつかのそのような実装によれば、整形ニューラルネットワーク520は、モジュール810に提供される同じ8個のサンプルを80回繰り返して、モジュール815への入力のための640サンプルを生成してもよい。しかしながら、代替的な実装では、整形ニューラルネットワーク520は、他の方法に従って声モデリング・ニューラルネットワーク525に入力を提供してもよい。たとえば、入力オーディオ・データは、16kHz以外の周波数でサンプリングされてもよく、モジュール805、810、815は異なるフレーム・サイズに対して動作してもよい、などである。ある代替的な実装では、モジュール805は、フレーム当たり20サンプルを受領してもよく、モジュール815は以前の20サンプルを履歴として使ってもよい。モジュール805、810、815のいくつかの詳細な例は、図14を参照して後述する。

【0098】

この例では、C_tは、現在の目標話者について声モデリング・ニューラルネットワーク525によって出力される合成されたオーディオ・データを表わす。図8には示されていないが、多くの実装において、トレーニング・プロセスは、C_tを損失関数決定ブロックに提供し、損失関数決定ブロックから損失関数値を受領し、損失関数値に従って少なくとも声モデリング・ニューラルネットワーク525をトレーニングすることに関わる。いくつかの実装は、損失関数値に従って整形ニューラルネットワークをトレーニングすることにも関わる。

【0099】

図9は、声モデリング・ニューラルネットワークの例示的なブロックおよび発話生成プロセスの間に声モデリング・ニューラルネットワーク提供されうる入力の例を示している。この例では、S_t、P_s→t、Q_s、F_sが声モデリング・ニューラルネットワーク525に提供される。ここで、S_tは目標話者識別データを表わし、P_s→tは声モデリング・ニューラルネットワーク525によって生成された、前の予測された(源から目標へのスタイル転移された(source-to-target style-transferred))オーディオ信号を表わし、Q_sは源話者の声に対応する、入力の時間整列された音素シーケンスを表わし、F_sは源話者の声に対応する基本的な周波数輪郭データを表わす。

【0100】

図8を参照して上記されるように、声モデリング・ニューラルネットワーク525は、この例において、モジュール805、810、815を含み、そのそれぞれは異なる時間分解能で動作する。よって、異なるフレーム・サイズの入力データが、たとえば整形ニューラルネットワーク520(図9には示さず)によって声モデリング・ニューラルネットワーク525に入力されてもよい。

【0101】

いくつかの実装は、一つまたは複数の追加的なニューラルネットワークに関わってもよい。いくつかのそのような実装では、第三のニューラルネットワークは、トレーニング・プロセス中に入力オーディオ・データ(たとえば図4を参照して上記した「第一の発話」)を受領してもよい。トレーニング・プロセスは、第一の人物の発話に対応する第一の発話特性を決定し、エンコードされたオーディオ・データを出力するよう第三のニューラルネットワークをトレーニングすることに関わってもよい。

【0102】

いくつかのそのような例によれば、第三のニューラルネットワークは、オートエンコー

10

20

30

40

50

ダであってもよく、あるいはオートエンコーダを含んでいてもよい。オートエンコーダは、効率的な符号化の教師なし学習のために使用されうるニューラルネットワークである。一般に、オートエンコーダの目標は、典型的には次元削減の目的のために、一組のデータについての表現または「エンコード」を学習することである。

【0103】

図10は、オートエンコーダのブロックの例を示す。オートエンコーダ1005は、たとえば、図3を参照して上記した制御システム315のような制御システムによって実装されてもよい。オートエンコーダ1005はたとえば、一つまたは複数の非一時的記憶媒体に記憶されたソフトウェアに従って、制御システムの一つまたは複数のプロセッサによって実装されてもよい。図10に示される要素の数および型は単に例である。オートエンコーダ1005の他の実装は、より多数、より少数または異なる要素を含んでいてもよい。

10

【0104】

この例において、オートエンコーダ1005は、三層のニューロンをもつ再帰型ニューラルネットワーク（recurrent neural network、RNN）を含む。いくつかの例によれば、オートエンコーダ1005のニューロンは、シグモイド活性化関数および／またはtanh活性化関数を用いてもよい。RNN層1～3におけるニューロンは、N次元入力データを、そのN次元状態を維持しつつ処理する。層1010は、RNN層3の出力を受領して、プーリング・アルゴリズムを適用するよう構成される。プーリングは、非線形ダウンサンプリングの一つの形である。この例によれば、層1010は、RNN層3の出力をM個の重ならない部分または「部分領域」の集合に分割して、そのような各部分領域について最大値を出力するmaxプーリング関数を適用するよう構成される。

20

【0105】

図11は、オートエンコーダを含む発話合成器をトレーニングするプロセスについての例示的なブロックを示している。この例では、トレーニング・プロセスの大半の側面が、図8を参照して上記したようにして実装されてもよい。S_t、F_t、Q_t、P_tは図8を参照して上記されている。

【0106】

しかしながら、この例では、オートエンコーダ1005も、声モデリング・ニューラルネットワーク525に入力を提供する。この例によれば、オートエンコーダ1005は、トレーニング・プロセス中に目標話者から入力オーディオ・データC_tを受領し、声モデリング・ニューラルネットワーク525にZ_tを出力する。この実装において、Z_tは、入力オーディオ・データC_tに比べて次元が削減されたオーディオ・データを出力する。

30

【0107】

この例において、C_t'は、現在の目標話者について声モデリング・ニューラルネットワーク525によって出力される合成されたオーディオ・データを表わす。この実装において、トレーニング・プロセスは、C_t'および「確固とした真実」オーディオ・データ——この例では入力オーディオ・データC_tである——を損失関数決定ブロック535に提供し、損失関数決定ブロックからの損失関数値を受領し、損失関数値に従って少なくとも声モデリング・ニューラルネットワーク525をトレーニングすることに関わる。この例は、損失関数値に従ってオートエンコーダ1005をトレーニングすることにも関わる。無用な混雑を避けるため、図10は、損失関数値が損失関数決定ブロック535によって声モデリング・ニューラルネットワーク525およびオートエンコーダ1005に提供されることを示す矢印を含んでいない。いくつかの実装は、損失関数値に従って整形ニューラルネットワークをトレーニングすることにも関わる。すべての場合において、トレーニング・プロセスは、トレーニングされるニューラルネットワークの少なくとも一つの重みに対応する少なくとも一つの有体な記憶媒体位置の物理的な状態を変更することに関わる。

40

【0108】

いくつかのそのような例では、トレーニングは、第三のニューラルネットワークによ

50

て生成されたエンコードされたオーディオ・データが第一の人物の発話に対応するかどうかを判定するよう第四のニューラルネットワークをトレーニングすることにも関わってもよい。第四のニューラルネットワークは、本稿では「話者素性分類器」または単に「話者分類器」と称されてもよい。いくつかのそのような実装では、発話生成プロセスは、第三のニューラルネットワークによって、源話者の発話に対応するオーディオ・データを受領することに関わってもよい。源話者は、いくつかの事例では、本稿の他所で（たとえば図6の記述において）言及される「第二の人物」に対応してもよい。したがって、受領されたオーディオ・データは、本稿の他所で言及される「第二の人物の第二の発話に対応する第二のオーディオ・データ」に対応してもよい。

【0109】

いくつかのそのような例では、発話生成プロセスは、第三のニューラルネットワークによって、第二のオーディオ・データに対応する第二のエンコードされたオーディオ・データを生成することに関わってもよい。発話生成プロセスは、第四のニューラルネットワークによって、第二のエンコードされたオーディオ・データを受領することに関わってもよい。いくつかの例では、発話生成プロセスは、第四のニューラルネットワークが修正された第二のエンコードされたオーディオ・データが第一の人物の発話に対応すると判定するまで、逐次反復プロセスを介して、修正された第二のエンコードされたオーディオ・データを生成し、第四のニューラルネットワークが修正された第二のエンコードされたオーディオ・データが第一の人物の発話に対応すると判定した後、修正された第二のエンコードされたオーディオ・データを第二のニューラルネットワークに（たとえば声モデリング・ニューラルネットワーク525）提供することに関わってもよい。

【0110】

図12Aは、話者分類器をトレーニングするプロセスの間に使用されうるブロックの例を示している。この例では、話者分類器1205は、ニューラルネットワークの一つの型であり、オートエンコーダ1005からの入力および損失関数決定ブロック535からのフィードバックに従ってトレーニングされる。話者分類器1205の、より詳細な例が、図12Bに示されており、下記に記述される。この実装によれば、話者分類器1205がトレーニングされる時点において、オートエンコーダ1005はすでにトレーニングされており、オートエンコーダ1005の重みは固定されている。

【0111】

この例によれば、オートエンコーダ1005は、トレーニング・プロセスの間に目標話者から入力オーディオ・データ C_t を受領し、話者分類器1205に Z_t を出力する。この実装では、 Z_t は、入力オーディオ・データ C_t に比べて次元が削減されている、エンコードされたオーディオ・データを含む。

【0112】

この実装によれば、話者分類器1205は、目標話者についての予測された話者識別データである S_{t^h} を、損失関数決定ブロック535に対して出力する。この例において目標話者についての「確固とした真実」話者識別データ S_t も損失関数決定ブロック535に入力される。

【0113】

「話者識別データ」に含まれるデータの型および量は、具体的な実装によって変わりうる。単純な場合には、話者識別データは、単に、特定の話者（たとえば「話者A」）が誰であることを示してもよい。いくつかのそのような事例では、話者分類器1205は、単に、たとえば話者が話者Aであるか、話者Aでないかを判定するようトレーニングされてもよい。いくつかのそのような実装によれば、話者分類器1205は、たとえば、損失関数決定ブロック535が S_{t^h} が S_t に一致すると判定するまでトレーニングされてもよい。

【0114】

しかしながら、いくつかの実装では、話者識別データは、より複雑であってもよい。いくつかのそのような実装によれば、話者識別データは、話者分類器1205によっておよび／またはオートエンコーダ1005によって学習された目標話者の発話特性を示しうる

10

20

30

40

50

。いくつかのそのような実装では、話者識別データは、目標話者の発話特性を表わす多次元ベクトルであってもよい。いくつかの実装では、ベクトルの次元は8、16、32、64または128であってもよい。いくつかのそのような実装によれば、話者分類器1205は、損失関数決定ブロック535が S_{t^h} と S_t の差が閾値以下であると判定するまで、トレーニングされてもよい。トレーニング・プロセスは、話者分類器1205の少なくとも一つの重みと対応する少なくとも一つの有体な記憶媒体位置の物理的な状態を変更することに関わる。

【0115】

図12Bは、話者分類器の一例を示している。この例では、話者分類器1205は、畳み込みニューラルネットワークを含む。この例によれば、話者分類器1205は、オートエンコーダ1005の出力を入力として受領し、この入力に基づいて話者分類を行なう。いくつかの例では、話者分類器1205への入力、 $M \times N$ の特徴を含む。ここで、 N は入力フレーム数、 M は特徴次元である。

10

【0116】

この例において、畳み込み層1210は64個のフィルタを含む。しかしながら、他の実装では、畳み込み層1210は30フィルタ、40フィルタ、50フィルタ、60フィルタ、70フィルタ、80フィルタ、90フィルタ、100フィルタなど異なる数のフィルタを含んでいてもよい。ここで、各フィルタ・カーネルは、 16×1 のフィルタ・サイズをもつ。この例によれば、畳み込み層1210は、ステップ・サイズ、または「ストライド」が4の畳み込み演算を実行する。ストライドは、フィルタをスライドさせるときに何個の特徴を通過するかを示す。よって、入力データに沿ったスライド・フィルタは、畳み込み層1210がこの例において実行する畳み込み演算の型である。この例における $M \times N$ 入力を与えられると、出力サイズは $C1 \times \text{floor}((N-16)/4+1)$ となる。ここで、 $\text{floor}(x)$ は $i \leq x$ となるような最大の整数 i を取る演算を表わす。

20

【0117】

この実装によれば、ニューラルネットワーク層1215は、畳み込み層1210からの出力を受領し、ReLU活性化関数を適用する。ここで、maxプール・ブロック1220はニューラルネットワーク層1215の出力に、maxプール演算を適用する。この例では、maxプール・ブロック1220は、8特徴毎から最大値を取ることによって、ニューラルネットワーク層1215の出力の次元を削減する。ここで、maxプール・ブロック1220は、 8×1 のカーネル・サイズをもち、ストライド8を適用する。

30

【0118】

この例において、畳み込み層1225は100個のフィルタを含む。しかしながら、他の実装では、畳み込み層1225は30フィルタ、40フィルタ、50フィルタ、60フィルタ、70フィルタ、80フィルタ、90フィルタ、110フィルタ、120フィルタ、130フィルタなど異なる数のフィルタを含んでいてもよい。ここで、各フィルタ・カーネルは、 5×1 のフィルタ・サイズをもつ。この例によれば、畳み込み層1225は、ステップ・サイズ、または「ストライド」が1の畳み込み演算を実行する。

【0119】

この実装によれば、ニューラルネットワーク層1230は、畳み込み層1225からの出力を受領し、ReLU活性化関数を適用する。ここで、maxプール・ブロック1235はニューラルネットワーク層1230の出力に、maxプール演算を適用する。この例では、maxプール・ブロック1235は、6特徴毎から最大値を取ることによって、ニューラルネットワーク層1230の出力の次元を削減する。ここで、maxプール・ブロック1220は、 6×1 のカーネル・サイズをもち、ストライド1を適用する。

40

【0120】

この例では、線形層1240は、maxプール・ブロック1235の出力を受領し、行列乗算を通じて線形変換を適用する。一つのそのような例によれば、線形層1240は

$$y = Ax + b \quad (\text{式1})$$

による線形変換を適用する。

50

【0121】

式1において、 x は入力を表わし、 A は学習可能な重み行列を表わし、 b は学習可能なバイアスを表わす。この実装によれば、層1245は線形層1240の出力にsoftmax関数を適用する。正規化指数関数としても知られるsoftmax関数は、任意の実数値の K 次元ベクトル z を、合計すると1になる範囲 $[0,1]$ 内の実数値の K 次元ベクトル $\sigma(z)$ に還元するロジスティック関数の一般化である。

【0122】

話者分類器1205の出力1250は、話者識別情報である。この例によれば、出力1250は、総数の話者素性クラスのうちの各話者素性クラスについて話者分類確率分布を含む。たとえば、出力1250は、話者が話者Aである確率 $P1$ 、話者が話者Bである確率 $P2$ などを含んでいてもよい。

10

【0123】

図13は、発話合成のために話者分類器およびオートエンコーダを使用することの例を与えている。このプロセスのいくつかの側面は、 C_s 、 S_t 、 $P_s \rightarrow t$ 、 F_s が何を表わすかも含め、上記してあり（たとえば図9を参照して）、ここでは繰り返さない。この例では、 C_s はオートエンコーダ1005に提供される源話者についての入力オーディオ・データを表わす。 Z_s は、話者分類器1205に輸入されるオートエンコーダ1005の出力を表わす。

【0124】

図13に示される時点では、オートエンコーダ1005および話者分類器1205は、トレーニング済みであり、その重みは記憶され、固定されている。この例によれば、損失関数決定ブロック535からのフィードバックはオートエンコーダ1005または話者分類器1205の重みを変更するためには使用されず、その代わりに、話者分類器1205が修正された源話者の発話を目標話者の発話であると分類するようになるまで Z_s の値を修正するために使われる。

20

【0125】

Z_s の値を修正するプロセスは、いくつかの実装では、確率的勾配降下プロセスのような逆最適化プロセスを含んでいてもよい。一例では、確率的勾配降下プロセスは、次のモデル関数 F ：

$$y = F(x, w) \quad (\text{式2})$$

に基づいていてもよい。

30

【0126】

式2において、 F は、パラメータ w 、入力 x および出力 y をもつモデル関数を表わす。式2に基づき、損失関数 L を：

$$\text{Loss} = L(F(x, w), Y) \quad (\text{式3})$$

のように構築してもよい。

【0127】

式3において、 Y は確固とした真実のラベルを表わす。ニューラルネットワークをトレーニングする通常のプロセスにおいては、 w の値を更新しようとする。しかしながら、この例では、ニューラルネットワークはトレーニング済みであり、 w の値は固定されている。よって、この例では、逆最適化プロセスは、式3における L の値を最小化するために x の値を更新または最適化することに関わってもよい。擬似コードでは、このプロセスの一例は、次のように記述できる。次のプロセスを、 $L(F(x, w), Y)$ が適切な最小に達するまで繰り返す：

40

$$\text{For } i=1, 2, \dots, n, \text{ do } x = x - \eta \nabla L_i(x) \quad (\text{式4})$$

【0128】

式4において、 η は学習率を表わす。この例において、 x は、オートエンコーダ1005から受領されたエンコードされたオーディオ・データを表わす。

【0129】

図14は、一例による、整列ニューラルネットワークおよび声モデリング・ニューラル

50

ネットワークのブロックを示す。図14のブロックは、たとえば、図3に示され、上記で述べた制御システム315のような制御システムを介して実装されてもよい。いくつかのそのような例では、図14のブロックは、一つまたは複数の非一時的媒体に記憶されたソフトウェア命令に従って制御システムによって実施されてもよい。図14に示されるフレーム・サイズ、フレーム・レート、要素数および要素の型は単に例である。

【0130】

この例によれば、整形ニューラルネットワーク520は、双方向RNN 1415を含む。この例では、整形ニューラルネットワーク520は、複数の話者（この例では109の話者）の発話特性データが記憶されているブロック1407をも含む。この例によれば、発話特性データは、各話者について整形ニューラルネットワーク520をトレーニングするプロセスの間に整形ニューラルネットワーク520によって学習された目標話者の発話特性に対応する。いくつかの実装では、ブロック1407は、発話特性データが記憶されるメモリ位置へのポインタを含んでいてもよい。この例では、発話特性データは32次元ベクトルによって表わされているが、他の例では、発話特性データは、他の次元のベクトルによってなど、他の仕方で表わされてもよい。

10

【0131】

この例では、ブロック1405は、たとえばユーザーから受領された入力に従って選択された特定の目標話者を表わす。ブロック1405からの話者識別データは、ブロック1407に提供され、ブロック1407は話者特性データ1410を連結ブロック1414に提供する。この実装では、連結ブロック1414は音素特徴1412をも（たとえば、上記の内容抽出ブロック505のような内容抽出ブロックから）受領する。

20

【0132】

この例では、連結ブロック1414は、発話特性データ1410を音素特徴1412と連結して、出力を、双方向RNN 1415に提供するように構成される。双方向RNN 1415は、たとえば、（たとえば図7を参照して）上記したように機能するように構成されてもよい。

【0133】

この例では、声モデリング・ニューラルネットワーク525は、モジュール805、810、815を含み、そのそれぞれは異なる時間分解能で動作する。この例では、モジュール805は、モジュール810よりもフレーム当たり、より多くのサンプルを処理し、モジュール815はモジュール805よりもフレーム当たり、より多くのサンプルを処理する。いくつかのそのような例では、モジュール805は、フレーム当たりモジュール810の10倍のサンプルを処理し、モジュール815は、フレーム当たりモジュール810の80倍のサンプルを処理する。いくつかのそのような実装によれば、声モデリング・ニューラルネットワーク525は、上記のような、SampleRNNニューラルネットワークの修正されたバージョンを含んでいてもよい。しかしながら、これらは単に例である。他の実装では、声モデリング・ニューラルネットワーク525は、異なる数のモジュールを含んでいてもよく、および／またはモジュールは異なるフレーム・サイズを処理するように構成されてもよい。

30

【0134】

モジュール805、810、815がそれぞれ異なる時間分解能で動作するので、この例では、整形ニューラルネットワーク520は、異なるフレーム・サイズの出力を、モジュール805、810、815のそれぞれに提供する。この例によれば、整形ニューラルネットワーク520は、50フレームに対応する時間区間、たとえば整形ニューラルネットワーク520が50個のフレームを生成した時間区間の間に1024サンプルのサイズをもつ50フレームをモジュール805に提供する。ここで、整形ニューラルネットワーク520は、1024サンプルのサイズをもつ50フレームを反復テンソル・ブロック1430に提供し、を反復テンソル・ブロック1430が該50フレームを10回反復し、モジュール810に、1024サンプルのサイズをもつ500個のフレームを提供する。この例において、整形ニューラルネットワーク520は、1024サンプルのサイズをもつ50フレームを反復テンソル・ブ

40

50

ロック 1 4 4 5 に提供し、を反復テンソル・ブロック 1 4 4 5 が該 50 フレームを 80 回反復し、モジュール 8 1 5 に、1024 サンプルのサイズをもつフレーム 4000 個を提供する。

【0 1 3 5】

この例によれば、モジュール 8 0 5 は、4000 個のオーディオ・サンプルをもつ単一のフレームをそれぞれ 80 サンプルをもつ 50 個のフレームに再編するよう構成された再編ブロック 1 4 1 8 を含む。この例では、線形演算ブロック 1 4 2 0 は、再編ブロック 1 4 1 8 によって出力されたフレーム当たり 80 サンプルから、フレーム当たり 1024 個のサンプルが生成されるよう、各フレームに対して線形演算を実行するよう構成される。いくつかのそのような例では、線形演算ブロック 1 4 2 0 は、次式：

$$Y=X*W \quad (\text{式5})$$

に従った行列乗算を介して各フレームに対して線形演算を実行するよう構成される。

【0 1 3 6】

式5において、X は入力行列を表わし、それはこの例では 50 かける 80 の次元をもち、再編ブロック 1 4 1 8 の出力に対応する。式5では、Y は次元 50 かける 1024 の出力行列を表わし、W は次元 80 かける 1024 の行列を表わす。よって、線形演算ブロック 1 4 2 0 の出力は、双方向 RNN 1 4 1 5 の出力と同数のフレームおよび同じフレーム・サイズをもち、それにより、線形演算ブロック 1 4 2 0 の出力および双方向 RNN 1 4 1 5 の出力は、合計され、RNN 1 4 2 2 に提供される。

【0 1 3 7】

線形演算ブロック 1 4 2 5 は、フレーム毎の 1024 サンプルが、RNN 1 4 2 2 によって出力されたフレーム毎の 1024 サンプルから生成されるよう、各フレームに対して線形演算を実行するよう構成される。いくつかの例では、線形演算ブロック 1 4 2 5 は、行列乗算を介して各フレームに対して線形演算を実行するよう構成される。再編ブロック 1 4 2 7 は、この例では、50 フレームの 10240 個のオーディオ・サンプルを、それぞれ 1024 サンプルをもつ 500 個のフレームに再編するよう構成される。よって、再編ブロック 1 4 2 7 の出力は、反復テンソル・ブロック 1 4 3 0 の出力と合計されうる。

【0 1 3 8】

この例において、再編ブロック 1 4 3 2 は、4000 サンプルをもつ入力オーディオ・データの一つのフレームを、それぞれ 8 サンプルをもつ 500 個のフレームに再編するよう構成される。線形演算ブロック 1 4 2 5 は、再編ブロック 1 4 3 2 によって出力されるフレーム当たり 8 個のサンプルから 1024 個のサンプルが生成されるよう、各フレームに対して線形演算を実行するよう構成される。いくつかの例では、線形演算ブロック 1 4 3 5 は、行列乗算を介して各フレームに対して線形演算を実行するよう構成される。これらの演算後、線形演算ブロック 1 4 3 5 の出力は、再編ブロック 1 4 2 7 の出力および反復テンソル・ブロック 1 4 3 0 の出力と合計されうる。

【0 1 3 9】

この合計は、この例では RNN 1 4 3 7 に提供される。ここで、RNN 1 4 3 7 の出力は線形演算ブロック 1 4 4 0 に提供され、線形演算ブロック 1 4 4 0 はこの例では、RNN 1 4 3 7 によって出力されるフレーム当たり 1024 サンプルから、フレーム当たり 8192 サンプルが生成されるよう、各フレームに対して線形演算を実行するよう構成される。いくつかの例では、線形演算ブロック 1 4 4 0 は、行列乗算を介して各フレームに対して線形演算を実行するよう構成される。この例では、再編ブロック 1 4 4 2 は、それぞれ 8192 サンプルをもつデータの 500 フレームを、それぞれ 1024 サンプルをもつ 4000 フレームに再編するよう構成される。すると、再編ブロック 1 4 4 2 の出力は、反復テンソル・ブロック 1 4 4 5 の出力と同じ次元をもち、よって、反復テンソル・ブロック 1 4 4 5 の出力と合計されうる。

【0 1 4 0】

この例によれば、モジュール 8 1 5 は、最近傍ブロック 1 4 4 7 を含む。最近傍ブロック 1 4 4 7 は、前の 7 個のサンプルを、オーディオ・データの現在のサンプルと一緒に、線形演算ブロック 1 4 5 0 に提供するよう構成される。この実装では、線形演算ブロック

10

20

30

40

50

1 4 5 0 は、最近傍ブロック 1 4 4 7 によって出力されるフレーム当たり 8 個のサンプルからフレーム当たり 1024 個のサンプルが生成されるよう、各フレームに対して線形演算を実行するよう構成される。すると、線形演算ブロック 1 4 5 0 の出力は、反復テンソル・ブロック 1 4 4 5 の出力と同じ次元をもち、よって、反復テンソル・ブロック 1 4 4 5 の出力および再編ブロック 1 4 4 2 の出力と合計されうる。いくつかの代替的な実装では、ブロック 1 4 4 7 および 1 4 5 0 は、たとえば全 1024 個のフィルタをもち各フィルタのサイズが 8 かける 1 である単一の畳み込み層によって、置き換えられてもよい。8×1 のフィルタ・サイズでは、畳み込みフィルタは前の 7 個のサンプルおよび現在のサンプルに対して作用することができる。

【0 1 4 1】

10

結果として得られる合計は、線形演算ブロック 1 4 5 2 に提供される。この例では、線形演算ブロック 1 4 5 2 は、線形演算ブロック 1 4 5 7 および線形演算ブロック 1 4 6 2 をも含む多層パーセプトロンの一部であるよう構成される。線形演算ブロック 1 4 5 2 の出力は ReLU ブロック 1 4 5 5 に提供され、ReLU ブロック 1 4 5 5 はその出力を線形演算ブロック 1 4 5 7 に提供する。線形演算ブロック 1 4 5 7 の出力は ReLU ブロック 1 4 6 0 に提供され、ReLU ブロック 1 4 6 0 はその出力を線形演算ブロック 1 4 6 2 に提供する。この例において、線形演算ブロック 1 4 6 2 は、ReLU ブロック 1 4 6 0 によって出力されたフレーム当たり 1024 サンプルからフレーム当たり 256 サンプルが生成されるよう、各フレームに対して線形演算を実行するよう構成される。この実装では、オーディオ・データは 8 ビット・オーディオ・データであり、よって、フレーム当たり 256 サンプルは、入力オーディオ・データの可能なオーディオ・サンプル値の数に対応する。

20

【0 1 4 2】

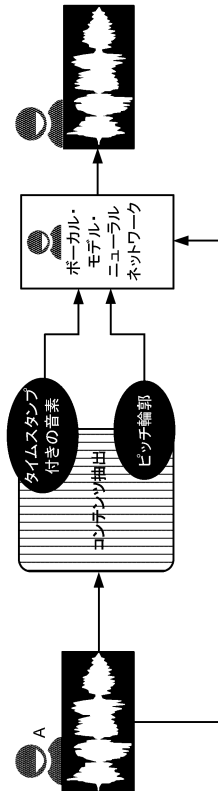
この実装によれば、ブロック 1 4 6 5 は線形演算ブロック 1 4 6 2 の出力に softmax 関数を適用する。この例では、softmax 関数は、256 個の値のそれぞれについての分類またはフレーム毎のクラスを提供する。この例では、声モデリング・ニューラルネットワーク 5 2 5 によって出力される出力データ 1 4 7 0 は、オーディオ・サンプル分布を含み、それが、256 個の値のそれぞれについての確率またはフレーム毎のクラスを示す。

【0 1 4 3】

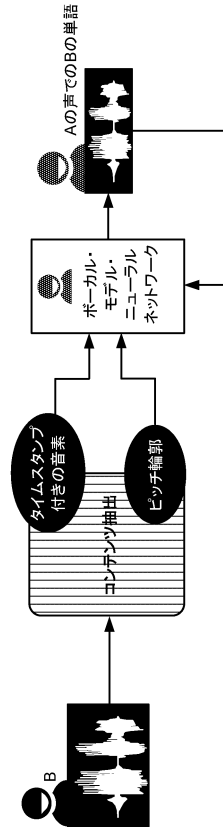
本稿で定義される一般原理は、付属の請求項の範囲から外れることなく他の実装に適用されてもよい。よって、請求項は、本稿に示される実装に限定されることは意図されておらず、本開示および本願で開示される原理および新規な特徴と整合する最も広い範囲を与えられるべきである。

30

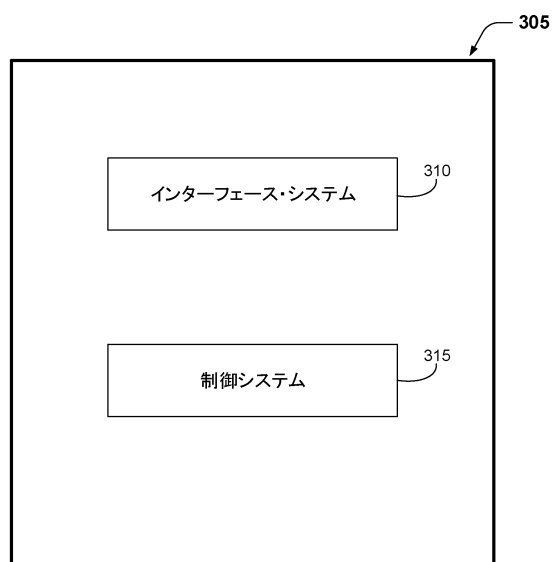
【図 1】



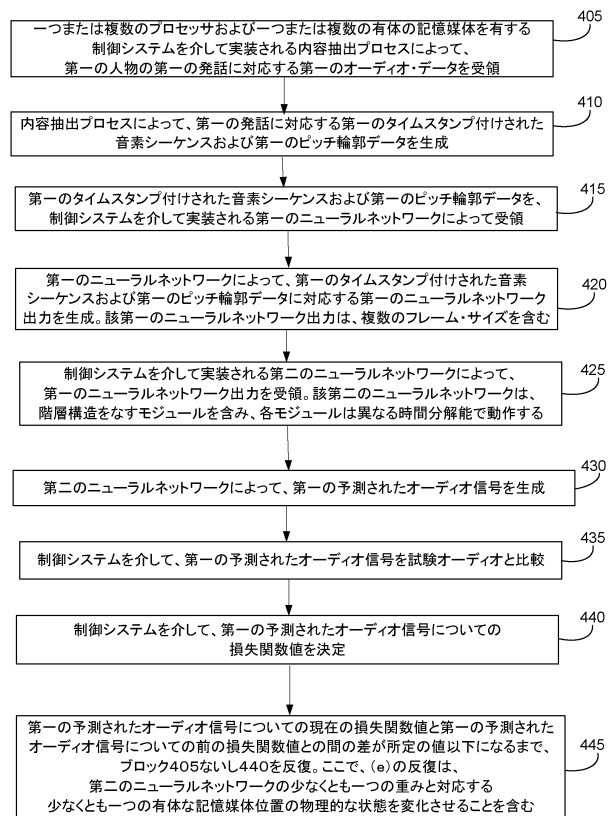
【図 2】



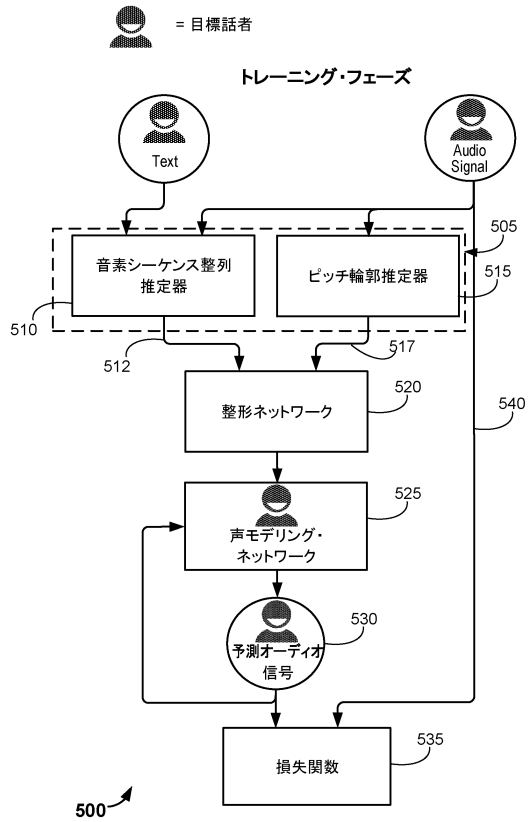
【図 3】



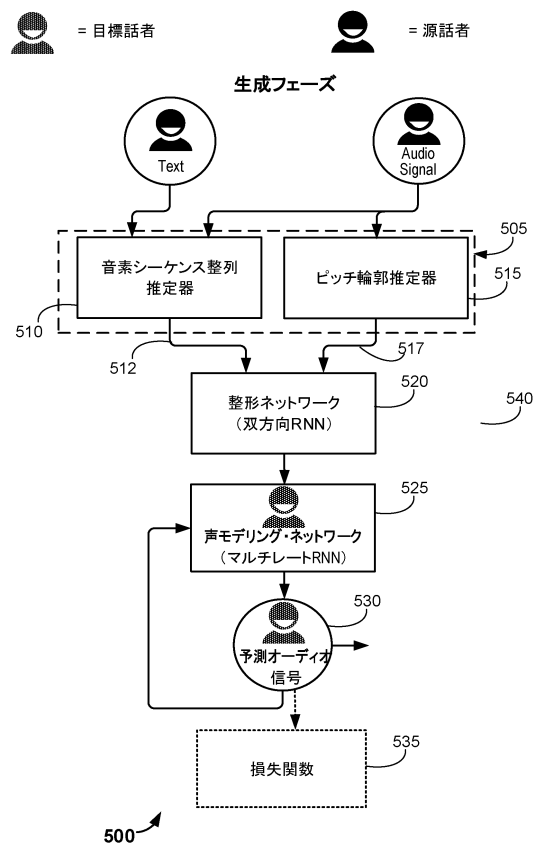
【図 4】



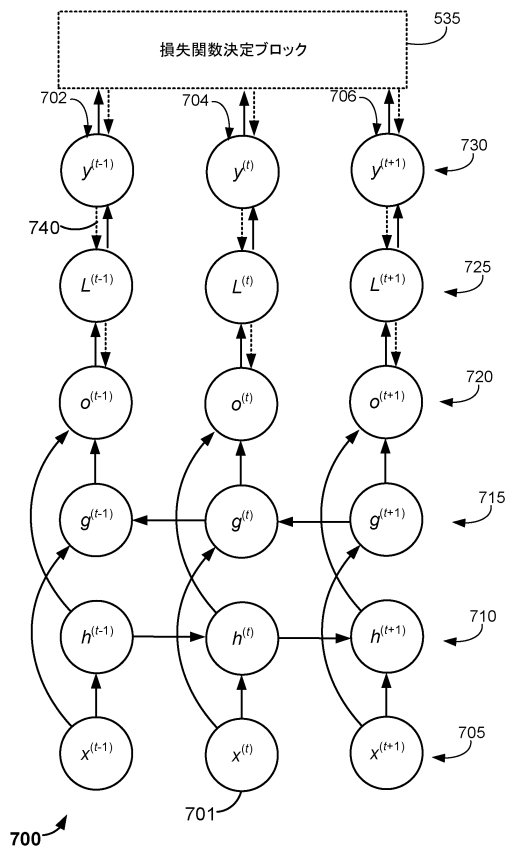
【図 5】



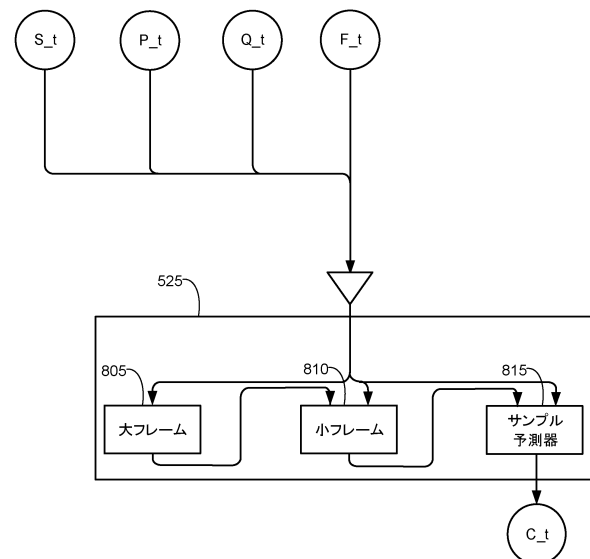
【図 6】



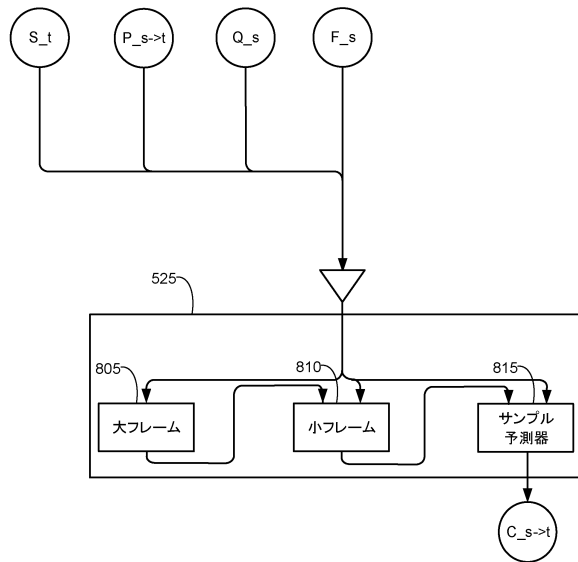
【図 7】



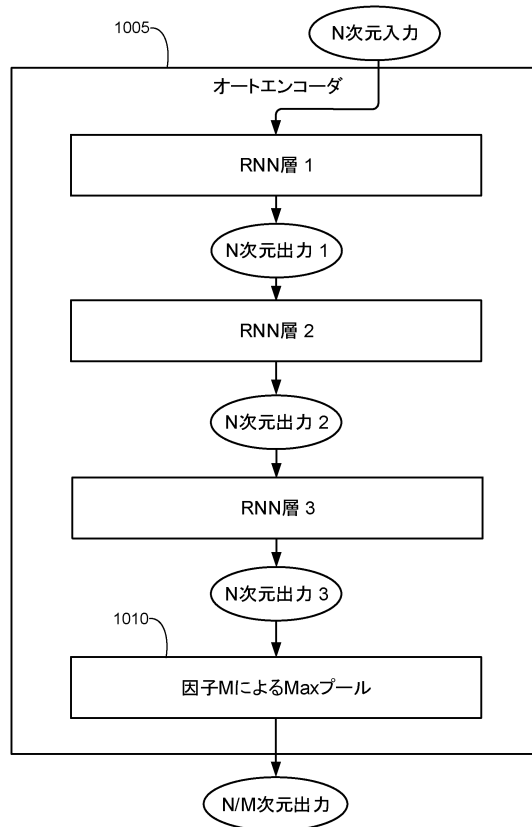
【図 8】



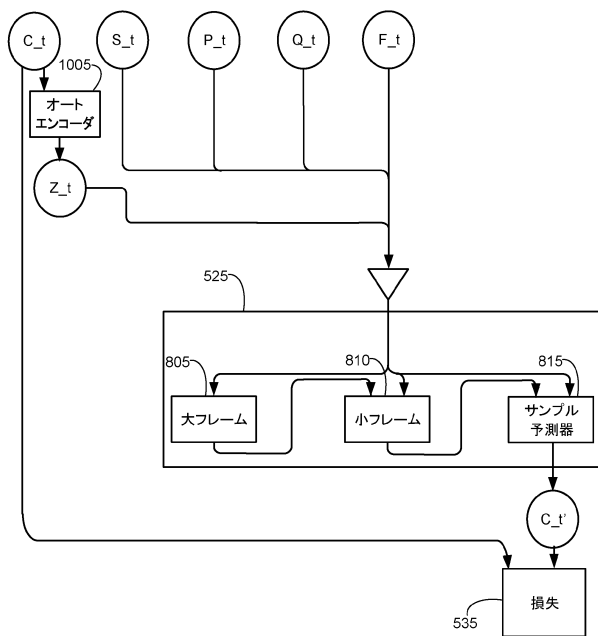
【図 9】



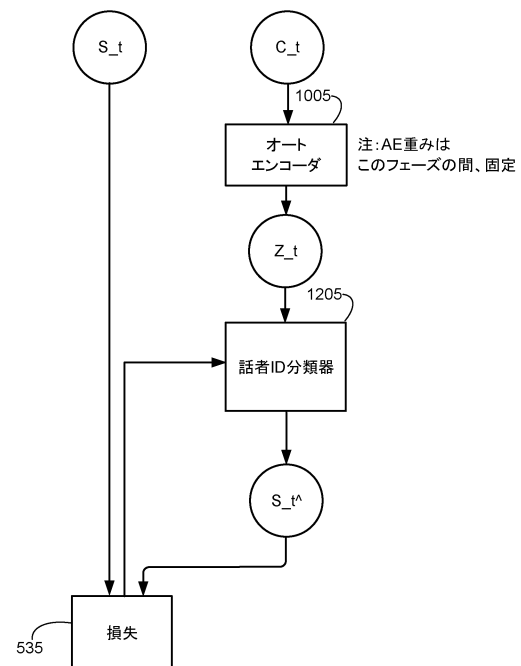
【図 10】



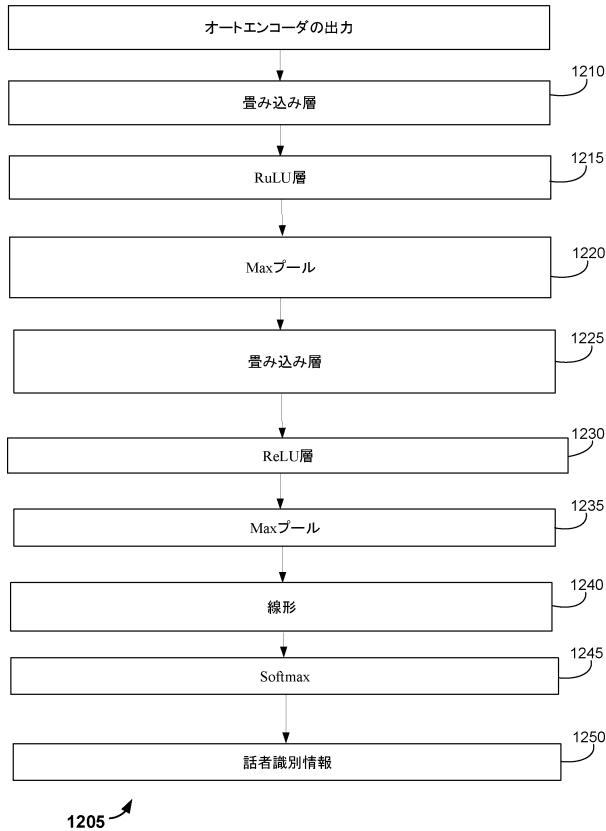
【図 11】



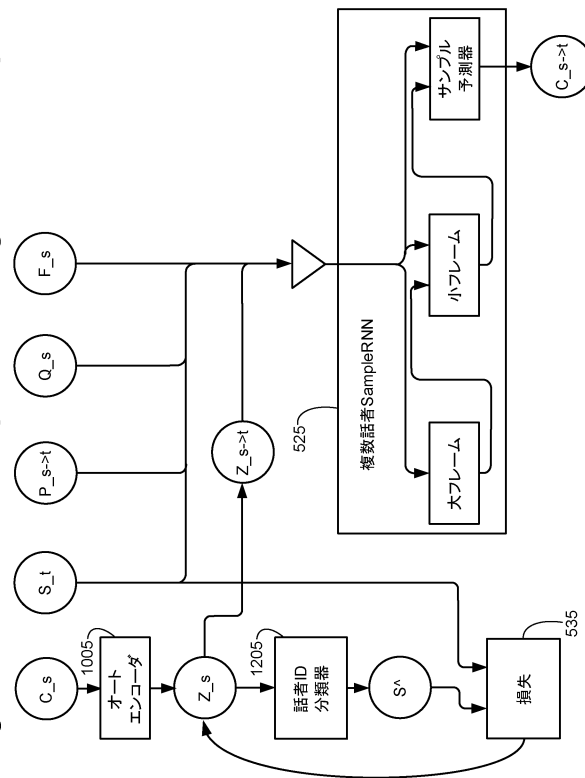
【図 12 A】



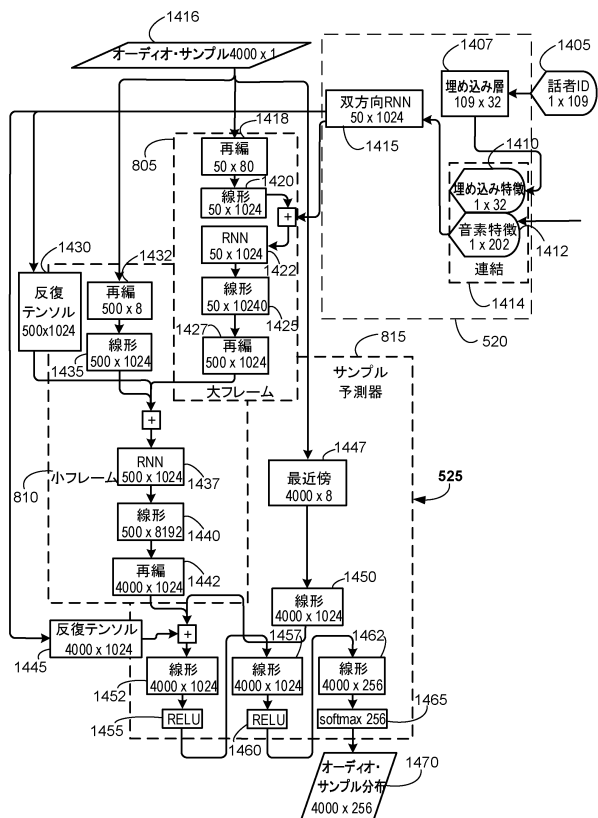
【図 1 2 B】



【図 1 3】



【図 1 4】



フロントページの続き

(31)優先権主張番号 62/797,864

(32)優先日 平成31年1月28日(2019.1.28)

(33)優先権主張国・地域又は機関
米国(US)

早期審査対象出願

(72)発明者 ジョウ, ツォーン

アメリカ合衆国 94103 カリフォルニア州 サンフランシスコ マーケット ストリート
1275 ドルビー ラボラトリーズ インコーポレイテッド 内

(72)発明者 ホーガン, マイケル ゲッティ

アメリカ合衆国 94103 カリフォルニア州 サンフランシスコ マーケット ストリート
1275 ドルビー ラボラトリーズ インコーポレイテッド 内

(72)発明者 クマール, ヴィヴェク

アメリカ合衆国 94103 カリフォルニア州 サンフランシスコ マーケット ストリート
1275 ドルビー ラボラトリーズ インコーポレイテッド 内

(72)発明者 モラレス, ジェイミー エイチ.

アメリカ合衆国 94501 カリフォルニア州 アラメダ ブロードウェイ 339 ユニット
ナンバー205

(72)発明者 ヴァスコ, クリスティーナ ミシェル

アメリカ合衆国 94112 カリフォルニア州 サンフランシスコ ミッション ストリート
3975 アpartment 1

審査官 上田 雄

(56)参考文献 米国特許出願公開第2018/0012613 (US, A1)

特開2017-003622 (JP, A)

米国特許出願公開第2016/0379622 (US, A1)

(58)調査した分野(Int.Cl., DB名)

G10L 21/00-25/93