

(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2004-301968

(P2004-301968A)

(43) 公開日 平成16年10月28日(2004.10.28)

(51) Int.Cl.<sup>7</sup>  
G10L 13/08F I  
G I O L 3/00テーマコード (参考)  
5 D O 4 5

審査請求 未請求 請求項の数 5 O L (全 10 頁)

(21) 出願番号 特願2003-92973 (P2003-92973)  
(22) 出願日 平成15年3月31日 (2003.3.31)(71) 出願人 000001487  
クラリオン株式会社  
東京都文京区白山5丁目35番2号  
(74) 代理人 100081961  
弁理士 木内 光春  
(72) 発明者 小原 優  
東京都文京区白山5丁目35番2号 クラ  
リオン株式会社内  
(72) 発明者 福永 功一郎  
東京都文京区白山5丁目35番2号 クラ  
リオン株式会社内

Fターム(参考) 5D045 AA09 AA20

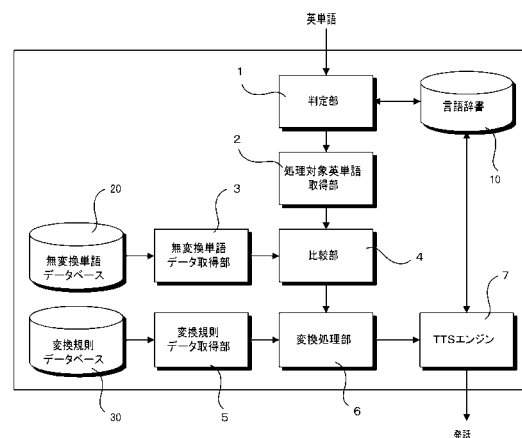
(54) 【発明の名称】 発話処理装置、発話処理方法及び発話処理用プログラム

## (57) 【要約】

【課題】 言語辞書に登録されていない単語を、ある程度正確に、もしくは理解できる単語に変換することができる発話処理装置を提供する。

【解決手段】 予め登録された単語の発音規則を保持する言語辞書と、変換すべきでない単語の一覧を保持する無変換単語データベースと、所定の変換規則を保持する変換規則データベースとを備え、発話対象となる単語が前記言語辞書に登録されているか否かを判断する判定部と、「変換処理の対象となる単語」を取得する処理対象単語取得部と、前記無変換単語データベースから「変換すべきでない単語」の一覧を取得する無変換単語データ取得部と、前記「変換処理の対象となる単語」と「変換すべきでない単語の一覧」とを比較して、変換すべき単語を選別する比較部と、前記変換すべき単語を、前記変換規則に従ってカタカナに変換する変換処理部と、この変換処理部によって変換された文字列を、前記言語辞書に保持された発音規則に従って音声化する音声合成部とを備える。

【選択図】 図1



## 【特許請求の範囲】

## 【請求項 1】

予め登録された単語の発音規則を保持する言語辞書と、  
変換すべきでない単語の一覧を保持する無変換単語データベースと、  
所定の変換規則を保持する変換規則データベースとを備えると共に、  
発話対象となる単語を取得して、この単語が前記言語辞書に登録されているか否かを判断する判定部と、  
その判定結果に基づいて、「変換処理の対象となる単語」を取得する処理対象単語取得部と、  
前記無変換単語データベースから、「変換すべきでない単語」の一覧を取得する無変換単語データ取得部と、  
前記「変換処理の対象となる単語」と「変換すべきでない単語の一覧」とを比較して、変換すべき単語を選別する比較部と、  
前記変換すべき単語を、前記変換規則に従ってカタカナに変換する変換処理部と、  
この変換処理部によって変換された文字列を、前記言語辞書に保持された発音規則に従って音声化する音声合成部と、  
を備えたことを特徴とする発話処理装置。

10

## 【請求項 2】

前記変換規則データベースが、アルファベットを保持している「アルファベットテーブル」と、あるアルファベットに付け加えられることで意味をなす文字列を保持している「続き文字テーブル」と、カタカナを保持している「カタカナテーブル」から構成されていることを特徴とする請求項 1 に記載の発話処理装置。

20

## 【請求項 3】

発話対象となる単語を取得して、その単語が、予め登録された単語についての発音規則を保持する言語辞書に登録されているか否かを判断する判定処理と、  
その判定結果に基づいて、「変換処理の対象となる単語」を取得する処理対象単語取得処理と、  
変換すべきでない単語の一覧を保持する無変換単語データベースから、その一覧を取得する無変換単語データ取得処理と、  
前記「変換処理の対象となる単語」と「変換すべきでない単語の一覧」とを比較して、変換すべき単語を選別する比較処理と、  
前記変換すべき単語を、取得された変換規則に従ってカタカナに変換する変換処理と、  
この変換処理によって変換された文字列を、前記言語辞書に保持された発音規則に従って音声化する音声合成処理と、  
を含むことを特徴とする発話処理方法。

30

## 【請求項 4】

前記変換規則データベースが、アルファベットを保持している「アルファベットテーブル」と、あるアルファベットに付け加えられることで意味をなす文字列を保持している「続き文字テーブル」と、カタカナを保持している「カタカナテーブル」から構成されていることを特徴とする請求項 3 に記載の発話処理方法。

40

## 【請求項 5】

コンピュータを制御することにより、発話対象となる単語を読み上げる発話処理用プログラムにおいて、  
そのプログラムは前記コンピュータに、  
発話対象となる単語を取得して、その単語が、予め登録された単語についての発音規則を保持する言語辞書に登録されているか否かを判断する判定ステップと、  
その判定結果に基づいて、「変換処理の対象となる単語」を取得する処理対象単語取得ステップと、  
変換すべきでない単語の一覧を保持する無変換単語データベースから、その一覧を取得する無変換単語データ取得ステップと、

50

前記「変換処理の対象となる単語」と「変換すべきでない単語の一覧」とを比較して、変換すべき単語を選別する比較ステップと、

前記変換すべき単語を、取得された変換規則に従ってカタカナに変換する変換ステップと、

この変換ステップにおいて変換された文字列を、前記言語辞書に保持された発音規則に従って音声化する音声合成ステップと、

を実行させるものであることを特徴とする発話処理用プログラム。

【発明の詳細な説明】

【0001】

【発明の属する技術分野】

TTS (Text To Speech) エンジンが英単語を読み上げる場合に参照する言語辞書に登録されていない英単語 (例えば、俗語やローマ字表記の名前) を、アルファベット読みではなく、日本語化した読み方で発話することができる発話処理装置、発話処理方法及び発話処理用プログラムに関するものである。

【0002】

【従来の技術】

従来から用いられている発話装置において、TTS エンジンが発話対象となる英単語を読み上げる場合、言語辞書よりその単語固有の発話規則を取得し、発話させるように構成されている。つまり、発話対象となる英単語を正確に発話するためには、言語辞書に多数の単語を登録しておく必要がある。この言語辞書は、容量を大きくすることで多くの発話規則をカバーできるが、メモリ容量等に制限がある場合には、正確に読み上げられる可能性は低くなる。

【0003】

このように言語辞書に登録されていない英単語の読み上げを要求された場合、その英単語を読み上げるための発話規則を取得することができないため、従来は、その英単語を単にアルファベット読みする方法がとられていた。また、ローマ字規則 (子音 + 母音、または母音のみで構成されている) に当てはまるアルファベットの羅列の場合には、その英単語が本来ローマ字読みすべきでないものであっても、それらをローマ字規則に従って日本語化して読み上げていた (特許文献 1、特許文献 2、特許文献 3 参照)。

【0004】

【特許文献 1】

特開 2002 - 23782 号公報

【特許文献 2】

特開平 11 - 305987 号公報

【特許文献 3】

特開 2000 - 10579 号公報

【0005】

【発明が解決しようとする課題】

しかしながら、上記言語辞書にも登録されておらず、聞き慣れない英単語をアルファベット読みされると、非常に聴き取り難い。また、発話スピードが早く設定されている場合には、なおさら聴き取り難くなる。さらに、アルファベットで読み上げられた英単語を頭の中で構築し直す必要があるため、集中力が阻害され、ナビゲーション装置等に用いた場合には、運転中等の動作に支障を与える可能性がある。

【0006】

また、上記言語辞書に登録されていない英単語を、単純にローマ字規則に沿って変換すると、その英単語がたまたまローマ字規則に当てはまるもの (例えば、more, amaze, are 等) であった場合には、一律にローマ字読みされてしまうため、正確に読み上げることができない。例えば、“more” は「モレ」、「amaze」は「アマゼ」、「are」は「アレ」と読み上げられてしまうことになる。

【0007】

10

20

30

40

50

本発明は、上述したような従来技術の問題点を解消するために提案されたものであり、その目的は、言語辞書に登録されていない英単語を、ある程度正確に、もしくは理解できる単語に変換することができる発話処理装置、発話処理方法及び発話処理用プログラムを提供することにある。

【 0 0 0 8 】

【課題を解決するための手段】

上記目的を達成するために、請求項 1 に記載の発話装置は、予め登録された単語の発音規則を保持する言語辞書と、変換すべきでない単語の一覧を保持する無変換単語データベースと、所定の変換規則を保持する変換規則データベースとを備えると共に、発話対象となる単語を取得して、この単語が前記言語辞書に登録されているか否かを判断する判定部と、その判定結果に基づいて、「変換処理の対象となる単語」を取得する処理対象単語取得部と、前記無変換単語データベースから、「変換すべきでない単語」の一覧を取得する無変換単語データ取得部と、前記「変換処理の対象となる単語」と「変換すべきでない単語の一覧」とを比較して、変換すべき単語を選別する比較部と、前記変換すべき単語を、前記変換規則に従ってカタカナに変換する変換処理部と、この変換処理部によって変換された文字列を、前記言語辞書に保持された発音規則に従って音声化する音声合成部とを備えたことを特徴とする。

10

【 0 0 0 9 】

また、請求項 3 に記載の発話処理方法は、請求項 1 に記載の発明を方法の観点で捉えたものであって、発話対象となる単語を取得して、その単語が、予め登録された単語についての発音規則を保持する言語辞書に登録されているか否かを判断する判定処理と、その判定結果に基づいて、「変換処理の対象となる単語」を取得する処理対象単語取得処理と、変換すべきでない単語の一覧を保持する無変換単語データベースから、その一覧を取得する無変換単語データ取得処理と、前記「変換処理の対象となる単語」と「変換すべきでない単語の一覧」とを比較して、変換すべき単語を選別する比較処理と、前記変換すべき単語を、取得された変換規則に従ってカタカナに変換する変換処理と、この変換処理によって変換された文字列を、前記言語辞書に保持された発音規則に従って音声化する音声合成処理とを含むことを特徴とする。

20

【 0 0 1 0 】

また、請求項 5 に記載の発明は、請求項 3 に記載の発明をコンピュータプログラムという観点で捉えたものであって、コンピュータを制御することにより、発話対象となる単語を読み上げる発話処理用プログラムにおいて、そのプログラムは前記コンピュータに、発話対象となる単語を取得して、その単語が、予め登録された単語についての発音規則を保持する言語辞書に登録されているか否かを判断する判定ステップと、その判定結果に基づいて、「変換処理の対象となる単語」を取得する処理対象単語取得ステップと、変換すべきでない単語の一覧を保持する無変換単語データベースから、その一覧を取得する無変換単語データ取得ステップと、前記「変換処理の対象となる単語」と「変換すべきでない単語の一覧」とを比較して、変換すべき単語を選別する比較ステップと、前記変換すべき単語を、取得された変換規則に従ってカタカナに変換する変換ステップと、この変換ステップにおいて変換された文字列を、前記言語辞書に保持された発音規則に従って音声化する音声合成ステップとを実行させるものであることを特徴とする。

30

40

【 0 0 1 1 】

上記のような構成を有する請求項 1 , 請求項 3 , 請求項 5 の発明によれば、音声合成部の保持する言語辞書に登録されていない単語の発話を、ある程度正確に行うことができる。また、言語辞書に登録されていない単語を、単純にアルファベット読みすることを防止できるので、聞き取りやすくなる。その結果、メールアドレス等のローマ字表記の文字列を、正確に読み上げることができるようになる。

【 0 0 1 2 】

請求項 2 に記載の発明は、請求項 1 に記載の発話処理装置において、前記変換規則データベースが、アルファベットを保持している「アルファベットテーブル」と、あるアルファ

50

ベットに付け加えられることで意味をなす文字列を保持している「続き文字テーブル」と、カタカナを保持している「カタカナテーブル」から構成されていることを特徴とする。

【 0 0 1 3 】

また、請求項 4 に記載の発明は、請求項 3 に記載の発話処理方法において、前記変換規則データベースが、アルファベットを保持している「アルファベットテーブル」と、あるアルファベットに付け加えられることで意味をなす文字列を保持している「続き文字テーブル」と、カタカナを保持している「カタカナテーブル」から構成されていることを特徴とする。

【 0 0 1 4 】

上記のような構成を有する請求項 2 又は請求項 4 の発明によれば、変換処理の対象となる単語の 1 文字目を抽出し、まず、「アルファベットテーブル」を参照することによりそのアルファベットを特定し、次に、あるアルファベットに付け加えられることで意味をなす文字列を保持している「続き文字テーブル」を参照して、2 文字目以降を特定し、最後に「カタカナテーブル」を参照して、カタカナ文字列に変換することができる。

【 0 0 1 5 】

【発明の実施の形態】

以下、本発明の実施の形態（以下、実施形態という）について、図面を参照して、具体的に説明する。

なお、本発明の各機能は、コンピュータを、ソフトウェアで制御することによって実現することが一般的である。この場合、コンピュータが備えるレジスタ、メモリ、外部記憶装置などの記憶装置が、いろいろな形式で、情報を一時的に保持したり永続的に保存する。そして、CPU が、前記ソフトウェアにしたがって、これらの情報に加工及び判断などの処理を加え、さらに、処理の順序を制御する。

【 0 0 1 6 】

また、コンピュータを制御するソフトウェアは、本出願の各請求項及び本明細書に記述する処理に対応した命令を組み合わせることによって作成され、作成されたソフトウェアは、アセンブルやコンパイルされた組み込みソフトウェアなどの形式で実行されることで、上記のようなハードウェア資源を活用する。

【 0 0 1 7 】

但し、本発明を実現するための上記のような態様はいろいろ変更することができ、例えば、本発明を実現するソフトウェアを記録した ROM チップや CD-ROM のような記録媒体は、それ単独でも本発明の一態様である。また、本発明の機能の一部を LSI などの物理的な電子回路で実現することも可能である。

【 0 0 1 8 】

（ 1 ）構成

図 1 は、本実施形態の発話処理装置の全体構成を示すブロック図である。すなわち、本実施形態の発話処理装置は、発話対象となる英単語を取得して、この英単語が後述する言語辞書 10 に登録されているか否かを判断する判定部 1 と、その判定結果に基づいて、変換処理の対象となる英単語を取得する処理対象英単語取得部 2 と、後述する無変換単語データベース 20 から、変換すべきでない単語の一覧を取得する無変換単語データ取得部 3 と、前記「変換処理の対象となる英単語」と「変換すべきでない単語の一覧」とを比較して、変換すべき英単語を選別する比較部 4 と、後述する変換規則データベース 30 から、変換規則データを取得する変換規則データ取得部 5 と、前記変換すべき英単語を、取得された変換規則に従ってカタカナに変換する変換処理部 6 と、この変換処理部 6 によって変換された文字列を、言語辞書 10 に保持された発音規則に従って音声化し、発話処理を行う TTS エンジン 7（請求項の音声合成部に相当）とを備えている。

【 0 0 1 9 】

また、前記言語辞書 10 は、日本語 / 英単語を問わず、単語の正確な発音規則を保持しているデータベースであり、この言語辞書 10 を参照して、前記 TTS エンジン 7 が発話処理を行うものである。

## 【 0 0 2 0 】

無変換単語データベース 20 は、“ m o r e ” , “ a r e ” 等、変換規則データベース 30 で間違った変換を行われてしまう可能性のある単語の一覧を保持しているデータベースである。すなわち、後述する変換規則データベース 30 に従うと、“ m o r e ” は「モレ」、 “ a r e ” は「アレ」と発話されてしまうが、このような誤った発話がなされる恐れがある単語を格納したものである。

## 【 0 0 2 1 】

また、変換規則データベース 30 は、ローマ字変換規則やフォニックス変換規則を保持するデータベースであり、例えば、以下の 3 種類のテーブルから構成されている。すなわち、アルファベットを保持している「アルファベットテーブル」と、あるアルファベットに付け加えられることで意味をなす文字列を保持している「続き文字テーブル」と、カタカナを保持している「日本語（カタカナ）テーブル」から構成されている。 10

## 【 0 0 2 2 】

なお、ローマ字変換規則とは、ローマ字のルールに沿った形式で、子音 + 母音 / 母音の組合せをカタカナにすることや、有効子音の範囲（ s h と s を同等にみなす等）を表すものであり、フォニックス変換規則とは、フォニックスに沿った形式で、子音 + 母音 / 母音の組合せをカタカナにするルールを表すものである。

## 【 0 0 2 3 】

なお、本発明に係る発話処理装置は、前記変換処理部 6 から、 T T S エンジン 7 へ文字列を受け渡す必要があるため、同一メモリ空間や、ネットワーク通信可能な範囲で動作させる必要がある。 20

## 【 0 0 2 4 】

## ( 2 ) 変換規則データベースの構成

上述したように、変換規則データベース 30 は、例えば 3 種類のテーブルから構成されているが、各テーブルの構成を以下に詳述する。

## 【 0 0 2 5 】

## ( 2 - 1 ) アルファベットテーブル

「アルファベットテーブル」は、アルファベットの全てを網羅しているテーブルであり、以下のように構成されている。

AlphTable[0] = "A" 30

AlphTable[1] = "B"

...

AlphTable[n] = "Z"

AlphTable[n+1] = NULL

## 【 0 0 2 6 】

## ( 2 - 2 ) 続き文字テーブル

続き文字テーブルは、以下のように構成されている。 40

```

Conv[0][0] = NULL
Conv[1][0] = "A"
Conv[1][1] = "YA"
Conv[1][2] = NULL
...
Conv[n][m + 1] = NULL
Conv[n+1][0] = NULL

```

10

## 【 0 0 2 7 】

このテーブルは、次のように用いられる。すなわち、上記「アルファベットテーブル」で `AlphTable[0]` は“ A ”なので、“ ア ”に変換される必要がある。そのため、 $n = 0$  の場合は、続き文字は存在しない。また、`AlphTable[1]` は“ B ”であるが、この“ B ”に“ A ”が続くと“ B A ”となり、“ バ ”になり得るため、続き文字は“ A ”となる。このことが、`Conv[1][0] = "A"` と表されている。また、`AlphTable[1]` の“ B ”は、“ バ ”以外にも、“ ビャ ”等があり得るため、続き文字として、“ Y A ”も準備しておく。

このように、続き文字テーブルは、各アルファベットに対して、変換可能な文字列を登録しておくものである。なお、`NULL`は、ターミネータ（終了宣言）である。 20

## 【 0 0 2 8 】

## ( 2 - 3 ) 日本語（カタカナ）テーブル

「日本語（カタカナ）テーブル」は、上記「アルファベットテーブル」と「続き文字テーブル」の配列番号から生成されるカタカナを格納しているテーブルであり、以下のように構成されている。なお、第 1 配列番号で  $n$  を使用し、第 2 配列番号で  $n + m$  としているのは、日本語テーブルを重複して参照されないようにするためである。

```

JapaneseTable[0][0] = "ア"
JapaneseTable[1][0] = "バ"
JapaneseTable[1][1] = "ビャ"
...

```

30

```

JapaneseTable[n][n+m] = "何かしらのカタカナ"

```

## 【 0 0 2 9 】

## ( 3 ) 作用

上記のような構成を有する本実施形態の発話処理装置における処理の流れを、図 2 に示したフローチャートを参照して説明する。

まず、判定部 1 が発話対象となる英単語を取得し（ステップ S 2 0 1）、この英単語が言語辞書 1 0 に登録されているか否かを判断する（ステップ S 2 0 2）。この英単語が言語辞書 1 0 に登録されていないと判断された場合には、ステップ S 2 0 3 に進み、この英単語を「変換処理対象となる英単語」として取得する。一方、ステップ S 2 0 2 において、発話対象となる英単語が言語辞書 1 0 に登録されていると判断された場合には、ステップ S 2 0 8 に進み、言語辞書 1 0 に保持されている発音規則に従って、TTS エンジン 7 によって発話される。 40

## 【 0 0 3 0 】

続いて、ステップ S 2 0 4 において、比較部 4 が、無変換単語データ取得部 3 を介して、無変換単語データベース 2 0 から「変換すべきでない単語の一覧」を取得し、ステップ S 2 0 3 において取得された「変換処理対象となる英単語」と「変換すべきでない単語の一 50

覧」とを比較して、変換すべき英単語を選別する（ステップS205）。

ステップS205において、「変換処理対象となる英単語」が、変換すべきでない単語であると判断された場合には、TTSエンジンが正確な読み上げ規則を保持しているので、そのままTTSエンジン7に渡され、発話される（ステップS208）。

#### 【0031】

一方、ステップS205において、「変換処理対象となる英単語」が変換すべき単語であると判断された場合には、変換処理部6が、変換規則データ取得部5を介して、変換規則データベース30から取得した変換規則データ（上記3種類のテーブル）に従って、変換すべき単語をカタカナに変換し（ステップS206、S207）、TTSエンジン7によって発話される（ステップS208）。

10

#### 【0032】

続いて、図2に示したフローチャートのステップ206～207の変換処理部6における処理について、図3のフローチャートを参照して説明する。

まず、図2のS205で取得された変換処理の対象となる単語（文字列）の1文字目を抽出する（ステップS301）。なお、このように、変換処理の対象となる単語を1文字ずつ分割することで、多言語への対応が可能となる。

#### 【0033】

続いて、図2のS206で、変換規則データベース30から取得した「アルファベットテーブル」を参照して、変換対象英単語の1文字目を特定することにより、nが検出される（ステップS302～ステップS303）。

20

#### 【0034】

次に、「続き文字テーブル」を参照して、変換対象英単語の2文字目以降の文字列が、続き文字テーブルで定義された文字列に当てはまるかどうかを確認する。その結果、mが検出される（ステップS304～ステップS305）。そして、特定されたn、mを用いて、日本語テーブルを参照し、カタカナ文字を取得する（ステップS306～ステップS307）。これを繰り返すことで、英単語を一文字ずつ、カタカナ文字列に変換することができる。このようにして変換された変換対象英単語をTTSエンジン7に渡し、発話される（図2のステップS208）。

#### 【0035】

##### （4）効果

上述したように、本実施形態の発話処理装置によれば、TTSエンジンの保持する言語辞書に登録されていない英単語の発話を、ある程度正確に行うことができる。また、言語辞書に登録されていない英単語を、単にアルファベット読みすることを防止できるので、聞き取りやすくなる。その結果、メールアドレス等のローマ字表記の文字列を、正確に読み上げることができるようになる。

30

#### 【0036】

##### （5）他の実施形態

本発明は、上述したような実施形態に限定されるものではなく、以下のような変形例が可能である。すなわち、図2のステップS205で「変換すべきでない単語」とされた英単語を、言語辞書10に自動的に登録するようにすれば、それ以後、その単語についてのステップS203～S205の処理が不要となる。また、変換規則データベースに格納される変換ルールは、適宜、追加・変更・削除が可能である。また、上記実施形態においては英単語を例に説明したが、他の言語に適用できることは言うまでもない。

40

#### 【0037】

##### 【発明の効果】

以上説明したように、本発明によれば、言語辞書に登録されていない英単語を、ある程度正確に、もしくは理解できる単語に変換することができる発話処理装置、発話処理方法及び発話処理用プログラムを提供することができる。

##### 【図面の簡単な説明】

【図1】本発明に係る発話処理装置の一実施形態の構成を示すブロック図

50

【図2】本発明の発話処理装置における処理の流れを示すフローチャート

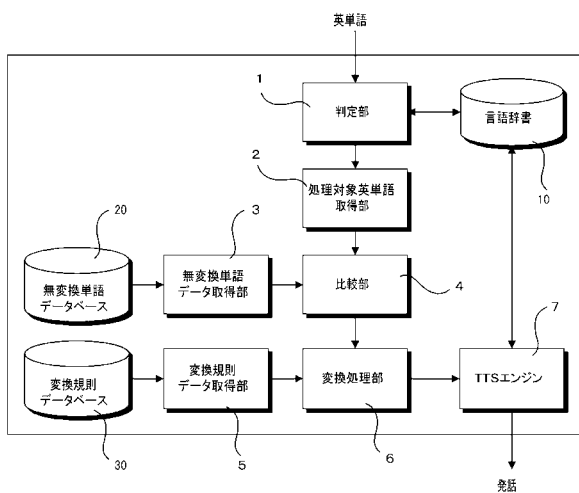
【図3】本発明の発話処理装置の変換処理部における処理の流れを示すフローチャート

【符号の説明】

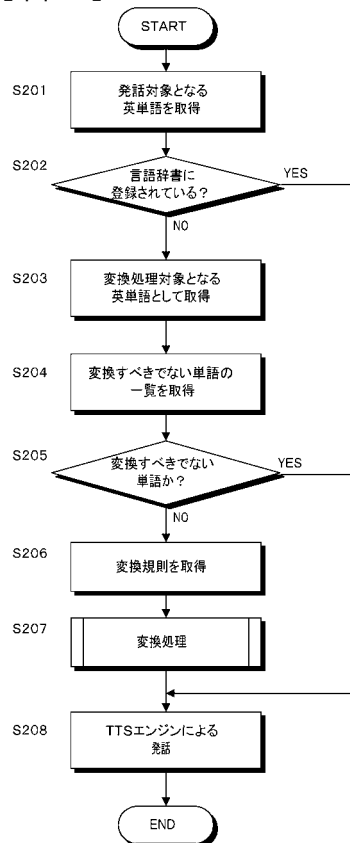
- 1 ... 判定部
- 2 ... 処理対象英単語取得部
- 3 ... 無変換単語データ取得部
- 4 ... 比較部
- 5 ... 変換規則データ取得部
- 6 ... 変換処理部
- 7 ... TTSエンジン
- 10 ... 言語辞書
- 20 ... 無変換単語データベース
- 30 ... 変換規則データベース

10

【図1】



【図2】



【図 3】

