

(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2007-232829

(P2007-232829A)

(43) 公開日 平成19年9月13日(2007.9.13)

(51) Int.Cl.	F I	テーマコード (参考)
G 1 0 L 13/00 (2006.01)	G 1 0 L 13/00 1 0 0 M	5 D 0 1 5
G 1 0 L 15/22 (2006.01)	G 1 0 L 15/22 2 0 0 H	
G 1 0 L 13/02 (2006.01)	G 1 0 L 13/02 1 2 1 A	
	G 1 0 L 13/02 1 3 0 B	

審査請求 有 請求項の数 5 O L (全 8 頁)

(21) 出願番号 特願2006-51771 (P2006-51771)
 (22) 出願日 平成18年2月28日 (2006.2.28)

(71) 出願人 000006297
 村田機械株式会社
 京都府京都市南区吉祥院南落合町 3 番地
 (74) 代理人 100086830
 弁理士 塩入 明
 (74) 代理人 100096046
 弁理士 塩入 みか
 (72) 発明者 新堂 安孝
 京都市伏見区竹田向代町 1 3 6 番地 村田
 機械株式会社内
 Fターム(参考) 5D015 AA01 AA04

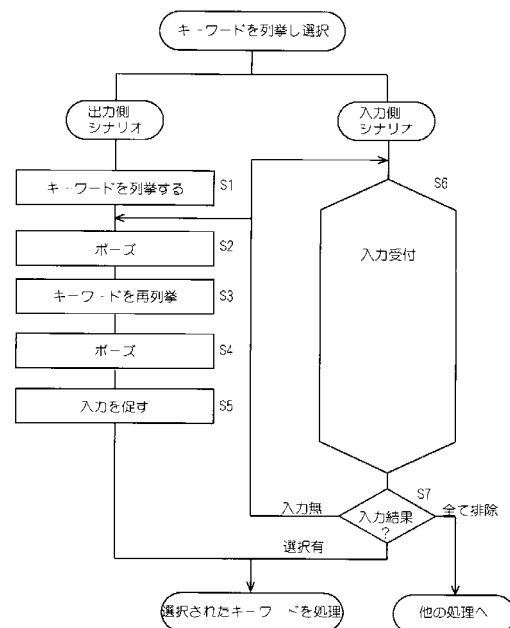
(54) 【発明の名称】 音声対話装置とその方法及びプログラム

(57) 【要約】

【課題】 音声対話装置が多数のキーワードを提示すると、人は個々のキーワードを聞き逃し易いし、また最適キーワードを選択するのが難しい。

【解決手段】 対話装置からキーワードを一旦列挙し、ポーズの後再度キーワードを提示して、選択を求める。有効な選択が有れば、選択肢に従ってシナリオを進め、選択がなければ再度キーワードを提示する。またどのキーワードも否定された場合、シナリオの他の場面へ分岐する。

【選択図】 図 4



【特許請求の範囲】**【請求項 1】**

人が音声入力するためのマイクロフォンと、マイクロフォンへの入力音声を音声認識するための音声認識装置と、スピーカを備えた音声出力装置と、シナリオの記憶部と、前記音声認識装置と音声出力装置とをシナリオに沿って制御する処理システム、とを備えた音声対話装置において、

複数のキーワードを列挙するようにスピーカから音声出力する際に、該複数のキーワードを最初に列挙した後に、音声出力をポーズし、次いで前記複数のキーワードを再度列挙して、人の入力音声を受け付けるように、前記記憶部に記憶するシナリオを構成することを特徴とする、音声対話装置。

10

【請求項 2】

再度の列挙では、最初に列挙したキーワードを、同順で、かつ少なくとも一部のキーワードを同義語に変換して音声出力するように、前記シナリオを構成することを特徴とする、請求項 1 の音声対話装置。

【請求項 3】

列挙されたキーワードに対する人からの入力音声を、遅くとも前記再度の列挙から処理するように、前記音声認識装置を構成することを特徴とする、請求項 1 の音声対話装置。

【請求項 4】

マイクロフォンで人の入力音声を受け付けて音声認識装置で音声認識し、予め記憶したシナリオに沿って、処理システムで音声認識装置と音声出力装置とを制御する音声対話方法において、

20

スピーカから複数のキーワードを列挙し、次いで音声出力をポーズし、さらに前記複数のキーワードを再度列挙して、人の入力音声を音声認識装置で認識する、ことを特徴とする、音声対話方法。

【請求項 5】

マイクロフォンで人の入力音声を受け付けて音声認識装置で音声認識し、予め記憶したシナリオに沿って、処理システムで音声認識装置と音声出力装置とを制御する音声対話用のプログラムにおいて、

スピーカから複数のキーワードを列挙する命令と、次いで音声出力をポーズする命令と、さらに前記複数のキーワードを再度列挙する命令と、少なくとも前記複数のキーワードを再度列挙する際に、人の入力音声を音声認識装置で認識する命令とを設けたことを特徴とする、音声対話プログラム。

30

【発明の詳細な説明】**【技術分野】****【0001】**

この発明は情報処理装置と人との音声による対話に関し、特にガイダンスなどのために装置側で予め記憶しているシナリオでの質問に人が答え易くすることに関する。

【背景技術】**【0002】**

音声対話装置が、人に多数のキーワードを列挙して選択を求める場面がある。このような場面で単純にキーワードを列挙されると、人は個々のキーワードを聞き逃し易い。そこでキーワード毎にポーズを入れると人に理解し易くなるが、各々のポーズ長の設定をどのようにするかなどで、シナリオの作成は複雑になる。

40

【特許文献 1】特開平 11 - 288292 号公報**【発明の開示】****【発明が解決しようとする課題】****【0003】**

この発明の課題は、多数のキーワードを音声で人に提示する際に、聞き取り易くかつ最適な選択をし易くすることにある。

請求項 2 の発明での追加の課題は、キーワードを繰り返すことにより対話が平板になり

50

、その結果冗長に感じられないように、しかも人が回答し易くすることにある。

請求項 3 , 5 の発明での追加の課題は、第 2 回目の列挙が終わるのを待たずに、人が回答できるようにすることにある。

【課題を解決するための手段】

【 0 0 0 4 】

この発明は、人が音声入力するためのマイクロフォンと、マイクロフォンへの入力音声を音声認識するための音声認識装置と、スピーカを備えた音声出力装置と、シナリオの記憶部と、前記音声認識装置と音声出力装置とをシナリオに沿って制御する処理システム、とを備えた音声対話装置において、複数のキーワードを列挙するようにスピーカから音声出力する際に、該複数のキーワードを最初に列挙した後に、音声出力をポーズし、次いで前記複数のキーワードを再度列挙して、人の入力音声を受け付けるように、前記記憶部に記憶するシナリオを構成することを特徴とする。

10

【 0 0 0 5 】

好ましくは、再度の列挙では、最初に列挙したキーワードを、同順で、かつ少なくとも一部のキーワードを同義語に変換して音声出力するように、前記シナリオを構成する。

また好ましくは、列挙されたキーワードに対する人からの入力音声を、遅くとも前記再度の列挙から処理するように、前記音声認識装置を構成する。

【 0 0 0 6 】

この発明は、マイクロフォンで人の入力音声を受け付けて音声認識装置で音声認識し、予め記憶したシナリオに沿って、処理システムで音声認識装置と音声出力装置とを制御する音声対話方法において、スピーカから複数のキーワードを列挙し、次いで音声出力をポーズし、さらに前記複数のキーワードを再度列挙して、人の入力音声を音声認識装置で認識する、ことを特徴とする。

20

人の入力音声の認識は例えば 2 回目の列挙から開始し、好ましくは 1 回目の列挙中、または列挙後のポーズから開始する。

【 0 0 0 7 】

この発明は、マイクロフォンで人の入力音声を受け付けて音声認識装置で音声認識し、予め記憶したシナリオに沿って、処理システムで音声認識装置と音声出力装置とを制御する音声対話プログラムにおいて、スピーカから複数のキーワードを列挙する命令と、次いで音声出力をポーズする命令と、さらに前記複数のキーワードを再度列挙する命令と、少なくとも前記複数のキーワードを再度列挙する際に、人の入力音声を音声認識装置で認識する命令とを設けたことを特徴とする。

30

【 0 0 0 8 】

この明細書において、音声対話装置に関する記載は音声対話方法や音声対話プログラムにもそのまま当てはまり、逆に音声対話方法に関する記載は音声対話装置や音声対話プログラムにもそのまま当てはまる。

キーワードの列挙に対する人の回答は例えばその選択である。

【発明の効果】

【 0 0 0 9 】

この発明では複数のキーワードを列挙して回答を求める際に、最初にキーワードを列挙し、一旦音声出力を停止して間を置き、次いでキーワードを再列挙する。このため最初の列挙でキーワードを聞き逃しても、次の列挙で正しく聞いて応答できる。最初の列挙と次の列挙の間にポーズが有るので、次の列挙が始まったときに、同じキーワードを繰り返していると直ちに分かる。さらに最初の列挙でキーワードの集合の概要を理解してから、キーワードが再度出力された時に回答すれば良く、回答が容易である。シナリオ作成の面では、ポーズ長を使い分ける必要が無く、ポーズの設定が簡単である。

40

【 0 0 1 0 】

2 回目の列挙でキーワードを同義語に変換して出力すると、対話が平板になり難い。そしてキーワードの語順を 1 回目と同じにすると、1 回目で人が聞いた順序でキーワードが繰り返されるので、回答し易い。

50

2 回目にキーワードを列挙する際には、人はほぼ回答できる状態にあるので、ここでキーワードを出力しながら回答を音声認識すると、人がキーワードを聞いて直ちに回答することに対応できる。

【発明を実施するための最良の形態】

【0011】

以下に本発明を実施するための最適実施例を示す。

【実施例】

【0012】

図1～図6に、実施例を示す。各図において、2は音声対話装置で、4は音声入力用のマイクロフォンで、6はアンプで省略可能である。8は音声認識装置で、10は音声認識に用いる辞書で、複数の辞書が対話装置2に記憶されている。12は音声の認識結果を出力するためのレジスタで、14は処理システムであり、16は音声対話用のシナリオ記憶部である。なおシナリオの個々のシーンを場面と呼び、各場面の記憶位置を番地と呼ぶ。

【0013】

シナリオ記憶部16に記憶するシナリオの構成を図2に示す。一般的シナリオ18はキーワードの列挙以外の部分のシナリオであり、キーワード列挙シナリオ20がキーワードの列挙に関する部分のシナリオである。なおキーワードの列挙とは2個以上のキーワードを列挙することで、好ましくは3個以上のキーワードを列挙することである。第1回キーワード列挙場面21で最初にキーワードを列挙し、ポーズ場面22で一旦ポーズし、第2回キーワード列挙場面23で再度キーワードを列挙する。ポーズ場面24で2回目のキーワードの列挙後にポーズし、入力促し場面25で2回目のポーズの後に入力を人に促す。シナリオには音声出力側のシナリオと音声入力側のシナリオとが同期して進行し、場面21～25は出力側のシナリオである。音声入力側のシナリオでは、入力受け付け場面26で列挙されたキーワードに対する人の選択を受け付けて音声認識する。そして列挙されたキーワードに対する音声認識は、例えば第1回キーワード列挙場面21、もしくはポーズ場面22、あるいは第2回キーワード列挙場面23から開始し、これに応じて図1の辞書10を列挙したキーワードに対応したものに切り替え、認識開始前にレジスタ12をゼロクリアする。

【0014】

図3にレジスタ12の構造を示す。ここではA～Gの7つのキーワードが列挙され、人に対してキーワードの選択を求めているものとする。有効な回答は、1以上のキーワードを選択することと、「全て要らない」、「要りません」などの全ての選択肢の否定である。レジスタ12には質問IDが記載され、これは入力側シナリオでの場面を示している。次の1ビットは回答が肯定か否定かを示し、肯定で0、否定でFである。各キーワードには同義語があり、例えば大学の学部案内の場合、「工学部」と「工」並びに「工学」は同義語である。そこで同義語をまとめて抽象化したものをサブジェクトと呼ぶと、人の回答はサブジェクトの選択と言い換えることができる。図3のレジスタ12では、A～Gの7個のサブジェクトに対し各1ビットを割り当てる。サブジェクトの個数は質問毎に変化するので、肯定/否定のビットの後のビットから、各サブジェクト毎に1ビット割り当てることとし、レジスタ12には充分大きな記憶容量を持たせる。

【0015】

「Aは要らないが、Bは要る」などのように肯定と否定が組み合わさった回答に備え、レジスタ12を複数段設け、「Aは要らない」を初段のレジスタで処理し、「Bは要る」を次段のレジスタで処理しても良い。また肯定/否定やサブジェクトの選択を各1ビットのデータで記憶する必要はなく、よりビット長の長いデータで記憶しても良い。

【0016】

辞書10には列挙する各キーワードとその同義語、並びにキーワードの範囲や組み合わせを示す単語、及び肯定/否定を表す語が記憶されている。キーワードの範囲や組み合わせを示す単語として「全部」、「全て」などがあり、大学の学部案内での理学と工学の組み合わせを示す単語として「理工学」、文学部と経済学部及び商学部の組み合わせを示す

ものとして「文化系」などがある。これらのキーワードや同義語等は、入力側シナリオの場面毎に辞書を変更して切り替える。肯定を表す語として「はい」や「お願い」、接尾辞の「たい」「です」「ます」などがある。否定を表す語として「いいえ」や接尾辞の「ない」などがある。肯定を表す語も否定を表す語も入力されない場合、肯定／否定のビットは初期値の肯定のままである。

【 0 0 1 7 】

音声認識装置 8 は入力音声の中に辞書 1 0 に記述された語があると、レジスタ 1 2 のこれに対応するビットへの書き込みを行う。単語が肯定あるいは否定を表す場合、肯定／否定のビットに対して 0 もしくは F を出力し、サブジェクト毎のビットでは該当ビットを F にする。またサブジェクトの集合に対応するキーワードを発見すると、その集合に含まれるサブジェクトのビットを F にセットする。そして音声認識装置 8 はキーワードを発見する都度、レジスタ 1 2 へ O R 加算により書き込みする。例えば大学の学部案内で「文学をお願い」との回答が入力された場合、「文学」をキーワードとして検出し、これに対するサブジェクトのビットを F にセットし、他のビットは 0 のままである。次に「お願い」が肯定に対応するので、先頭の肯定ビットを 0 に保ち、他のビットの値を変えない。この場合のレジスタ 1 2 の出力は、肯定ビットが 0 で肯定であり、「文学」のビットがセットされ、他はセットされていないので、文学についてのみ案内を依頼したことになる。「文学と経済をお願い」の場合、「文学」のビットと「経済」のビットがセットされ、肯定／否定のビットは肯定の 0 ビットに保たれる。

10

【 0 0 1 8 】

列挙されたキーワードに対する選択を認識する際の特別のルールとして、「はい」「それ」などのキーワードを特定しない入力には、その直前に出力されたキーワードが肯定で選択されたものとする。これは 2 回目の列挙の途中で「はい」などが入力されることに備えたものであるが、このルールは設けなくても良い。また「文学は要らないが、経済学を知りたい」などの肯定／否定の述語が 2 語以上ある入力に対し、レジスタ 1 2 を複数段設けると、1 段目のレジスタで「文学は要らない」に対し肯定／否定のビットを否定の F にセットし、「文学」のビットを F にセットする。次の段のレジスタで「経済学を知りたい」を処理し、肯定／否定のビットを肯定の 0 にセットし、「経済学」のビットを F にセットする。この場合の認識結果は、「経済学部について知りたい」というものと同じである。

20

30

【 0 0 1 9 】

図 1 に戻り、音声データ生成部 3 0 はシナリオに基づいて音声を生成し、アンプ 3 2 を介してスピーカ 3 4 から音声出力する。なおアンプ 3 2 は設けなくても良い。実施例では音声対話装置 2 をガイダンス用のロボットに組み込み、処理システム 1 4 からのジェスチャー信号でロボット本体 3 6 を動作させる。

【 0 0 2 0 】

図 4 に、実施例の音声対話方法を示す。シナリオ中のキーワードを列挙し選択する部分で、出力側のシナリオではステップ 1 でキーワードを列挙する。次いでステップ 2 でポーズし、ステップ 3 でキーワードを再列挙する。そしてステップ 4 でポーズした後、ステップ 5 で入力を促す。ただしステップ 4 のポーズは省略しても良い。実施例の場合、ステップ 2 あるいはステップ 4 で、ロボット本体 3 6 に身振りをさせても良い。またステップ 3 でキーワードを再列挙する場合、ステップ 1 で列挙したキーワードの一部を同義語に変換し、特によりシンプルな単語に変更して同じ順で列挙すると、1 回目と 2 回目とで表現を変えてしかも同じ順なので、人が答え易くでき、しかも冗長さを弱めることができる。

40

【 0 0 2 1 】

入力側シナリオでは、ステップ 1 のキーワードの列挙から、入力を受け付け、入力音声を音声認識する。ただしステップ 2 のポーズから、もしくはステップ 3 のキーワードの再列挙から、入力音声を音声認識しても良い。そしてステップ 7 で入力結果を判別し、有効な入力がなかった場合、ステップ 2 のポーズまでもしくはステップ 3 のキーワードの再列挙まで戻って、あるいはキーワードの再列挙に準じた処理を行って、入力を再度受け付け

50

る。全ての選択肢が否定された場合他の処理へ移り、いずれかのキーワードが選択されると、選択されたキーワードやその組み合わせについてガイダンスする。

【 0 0 2 2 】

図 5 に音声対話プログラム 4 0 の構造を示すと、一般シナリオ命令 4 1 はキーワードの列挙以外のシナリオを処理し、キーワード列挙命令 4 2 は 1 回目のキーワードの列挙を処理し、キーワード再列挙命令 4 3 は 2 回目のキーワードの列挙を処理する。ポーズ命令 4 4 は、1 回目のキーワードの列挙と 2 回目のキーワードの列挙の間、並びに 2 回目のキーワードの列挙後の、ポーズやジェスチャーを処理する。音声認識命令 4 5 は例えば 1 回目のキーワード列挙中から音声認識を開始し、図 1 の辞書 1 0 を列挙されたキーワードに合わせて切り替える。音声認識命令 4 5 は入力音声の認識結果に基づき、シナリオでのそれ
10
以前の処理へ戻る、他の処理へ移る、選択されたキーワードについてガイダンスする、ようにシナリオを分岐させる。促し命令 4 6 は、2 回目のキーワードの列挙後と 2 回目のポーズの後に、入力を促す文を出力する。

【 0 0 2 3 】

図 6 に、大学の学部案内を例に音声ガイダンスの具体例を示す。なおこの具体例は実施例の音声対話装置、音声対話方法、音声対話プログラムのいずれにも該当する。ステップ 1 1 で、大学の学部としてどのようなものがあるかを列挙する。ステップ 1 2 で、ロボット本体のジェスチャーを交えてポーズし、ステップ 1 3 でキーワードを再列挙する。ここで文学を「文」、経済学を「経済」としているように、キーワードをより短いキーワードで同義語のものに変換して同じ順序で列挙する。ステップ 1 4 で再度ポーズし、ステップ
20
1 5 で回答を促す文を出力する。

【 0 0 2 4 】

入力側シナリオでは、ステップ 1 1 の段階で「経済学」などの回答がされることに備えて、ステップ 1 1 のキーワードの列挙から入力音声を認識する。ただしステップ 1 2 のポーズから、もしくはステップ 1 3 の再列挙から、入力音声の認識を開始しても良い。そしてステップ 1 7 で、入力結果に基づき分岐する。

【 0 0 2 5 】

実施例では以下の効果がある。

- (1) キーワードが 2 回以上列挙されるので、人にとって聞き逃しにくい。
- (2) 1 回目のキーワードの列挙でキーワードの全体をおおよそ理解し、2 回目のキーワードの列挙で正確に聞き取って回答すればよいので、正確な回答が容易である。
30
- (3) 1 回目と 2 回目とでキーワードの列挙の仕方を変えるので、対話が平板になりにくい。
- (4) サブジェクト毎のビット和を用いるので、「文学」,「経済学」などの個別の回答から、「文化系」,「全部」などの範囲を示す回答まで対応できる。なお「A, B, C は要らない」などの入力に対し、これ以外のキーワードが選択されたものとして解釈すると、入力を認識できる範囲をさらに広げることができる。

【 図面の簡単な説明 】

【 0 0 2 6 】

【 図 1 】 実施例の音声対話装置のブロック図

【 図 2 】 実施例で用いるシナリオのブロック図

【 図 3 】 実施例で用いる音声認識用のレジスタを示す図

【 図 4 】 実施例の音声対話方法を示すフローチャート

【 図 5 】 実施例の音声対話プログラムのブロック図

【 図 6 】 実施例を大学の学部案内に応用した例を示す図

【 符号の説明 】

【 0 0 2 7 】

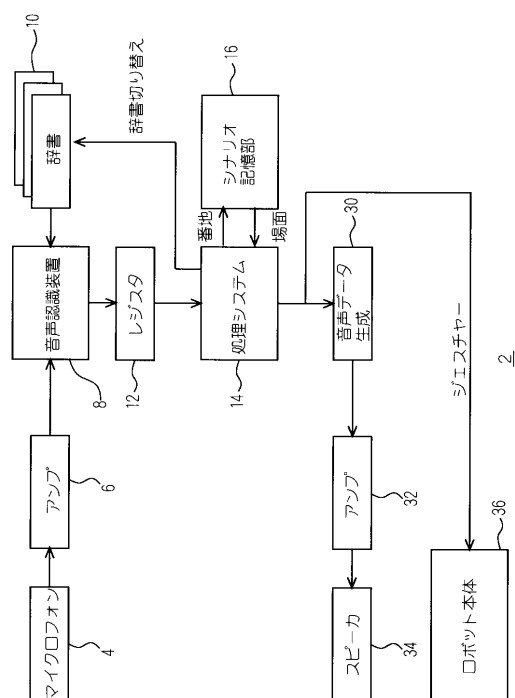
- 2 音声対話装置
- 4 マイクロフォン

6	3	2	アンプ
8			音声認識装置
1	0		辞書
1	2		レジスタ
1	4		処理システム
1	6		シナリオ記憶部
1	8		一般的シナリオ
2	0		キーワード列挙シナリオ
2	1		第1回キーワード列挙場面
2	2		ポーズ場面
2	3		第2回キーワード列挙場面
2	4		ポーズ場面
2	5		入力促し場面
2	6		入力受け付け場面
3	0		音声データ生成部
3	4		スピーカ
3	6		ロボット本体
4	0		音声対話プログラム
4	1		一般シナリオ命令
4	2		キーワード列挙命令
4	3		キーワード再列挙命令
4	4		ポーズ命令
4	5		音声認識命令
4	6		促し命令

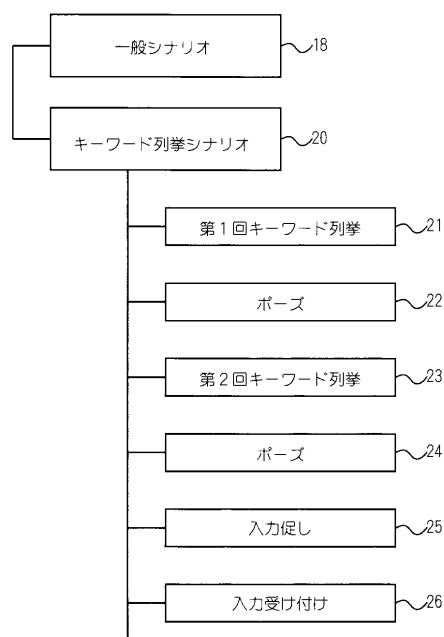
10

20

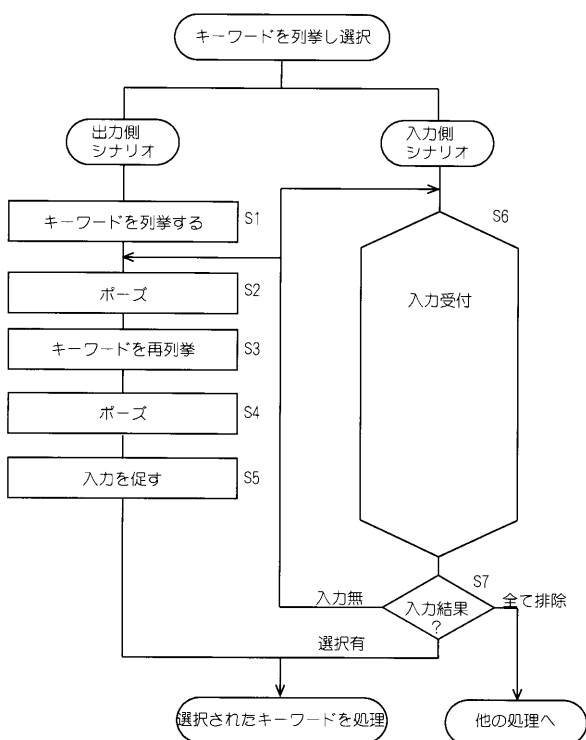
【 圖 1 】



【 圖 2 】



【 图 4 】



【图 6】

