

(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2011-237602

(P2011-237602A)

(43) 公開日 平成23年11月24日(2011.11.24)

(51) Int.Cl.	F I	テーマコード (参考)
G 1 O G 3/04 (2006.01)	G 1 O G 3/04	5 D O 8 2
G 1 O L 19/00 (2006.01)	G 1 O L 19/00 2 2 O Z	
G 1 O L 21/04 (2006.01)	G 1 O L 21/04 1 2 O C	
G 1 O L 13/02 (2006.01)	G 1 O L 13/02 1 2 2 B	

審査請求 未請求 請求項の数 22 O L (全 44 頁)

(21) 出願番号	特願2010-109006 (P2010-109006)	(71) 出願人	000002897
(22) 出願日	平成22年5月11日 (2010. 5. 11)		大日本印刷株式会社
			東京都新宿区市谷加賀町一丁目1番1号
		(74) 代理人	100091476
			弁理士 志村 浩
		(72) 発明者	茂出木 敏雄
			東京都新宿区市谷加賀町一丁目1番1号
			大日本印刷株式会社内
		Fターム(参考)	5D082 BB01 BB13 BB14 BB15

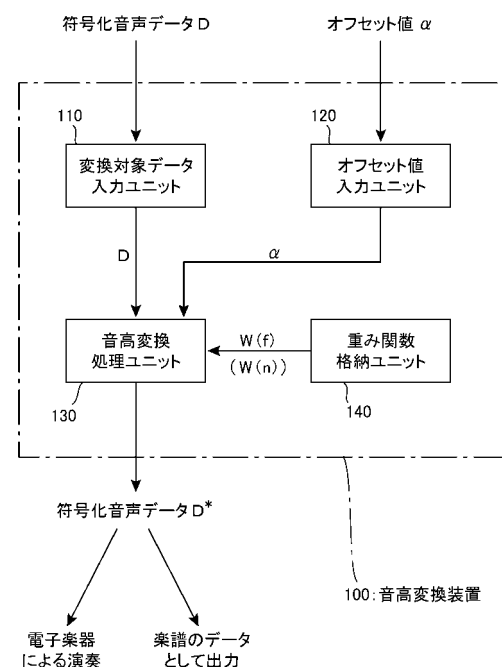
(54) 【発明の名称】 符号化音声データの音高変換装置

(57) 【要約】

【課題】人間の声を表現した符号化音声データDに対して、簡便な方法により音高シフトを施し、違和感のない自然な音声再生を可能にする。

【解決手段】ユニット110により変換対象となる符号化音声データDを入力し、ユニット120により音高のシフト量を示すオフセット値 α を入力する。ユニット140内には、周波数 f の増加に従って単調減少する重み関数 $W(f)$ を格納しておく。ユニット130は、データD内の各符号について、当該符号が示す周波数 f を、 $\alpha \cdot W(f)$ に応じた値だけ増減する処理を行い、処理後の符号を符号化音声データ D^* として出力する。各符号に対する音高の実際のシフト量は、重み関数 $W(f)$ を乗じて決定されるため、より高音の符号ほど音高シフト量が減少する。このため、フォルマント成分についての音高シフト量が小さくなり、簡便ながら、自然な音声再生が可能な符号化音声データ D^* が得られる。

【選択図】 図12



【特許請求の範囲】

【請求項 1】

特定周波数の音が特定時間だけ持続することを示す符号を時間軸上に並べることによって構成され人間の声を表現する符号化音声データを、異なる音高をもった別な音声データに変換する符号化音声データの音高変換装置であって、

変換対象となる符号化音声データ D を入力する変換対象データ入力ユニットと、

音高に関するオフセット値 α を入力するオフセット値入力ユニットと、

周波数 f について定義された所定の重み関数 $W(f)$ を格納した重み関数格納ユニットと、

前記変換対象となる符号化音声データ D に対して、前記オフセット値 α に基づく音高の変更処理を行い、変更後の符号化音声データ D^* を出力する音高変換処理ユニットと、
を備え、

前記重み関数格納ユニットが、前記重み関数 $W(f)$ として、周波数 f 軸上の所定区間において $W(f)$ が周波数 f の増加に従って単調減少する関数を格納しており、

前記音高変換処理ユニットが、前記重み関数 $W(f)$ を用いて、前記変換対象となる符号化音声データ D に含まれている個々の符号について、当該符号が示す周波数 f を $\alpha \cdot W(f)$ に応じた値だけ増減することにより新たな周波数 f' を求め、当該符号を、それが示す周波数 f を f' に変更した新たな符号に置き換える処理を行うことを特徴とする符号化音声データの音高変換装置。

【請求項 2】

請求項 1 に記載の音高変換装置において、

重み関数格納ユニットが、重み関数 $W(f)$ として、周波数 f 軸上の第 1 設定値 f_a および第 2 設定値 f_b ($f_b > f_a$) について、 $f \leq f_a$ の区間は、 $W(f) = 1$ 、 $f_a < f < f_b$ の区間は、 $1 > W(f) > 0$ (但し、 $W(f)$ は f の増加に従って単調減少)、 $f \geq f_b$ の区間は、 $W(f) = 0$ となる関数を格納していることを特徴とする符号化音声データの音高変換装置。

【請求項 3】

請求項 2 に記載の音高変換装置において、

重み関数格納ユニットが、 $100\text{Hz} \sim 200\text{Hz}$ の範囲内の第 1 設定値 f_a と $3\text{kHz} \sim 6\text{kHz}$ の範囲内の第 2 設定値 f_b とを用いた男性用重み関数 $W\text{-male}(f)$ と、 $200\text{Hz} \sim 400\text{Hz}$ の範囲内の第 1 設定値 f_a と $4\text{kHz} \sim 8\text{kHz}$ の範囲内の第 2 設定値 f_b とを用いた女性用重み関数 $W\text{-female}(f)$ と、を格納しており、

音高変換処理ユニットが、変換対象となる符号化音声データ D によって表現される声が男性の声か女性の声かを示す指示に基づいて、男性の声の場合には前記男性用重み関数 $W\text{-male}(f)$ を用いた変更処理を行い、女性の声の場合には前記女性用重み関数 $W\text{-female}(f)$ を用いた変更処理を行うことを特徴とする符号化音声データの音高変換装置。

【請求項 4】

請求項 2 または 3 に記載の音高変換装置において、

オフセット値入力ユニットが、音高を高める場合は正、低める場合は負のオフセット値 α を入力し、

音高変換処理ユニットが、所定の係数 k ($k > 1$) を用いた式 $f' = f \cdot k^{\alpha} \cdot W(f)$ により新たな周波数 f' を求めることを特徴とする符号化音声データの音高変換装置。

【請求項 5】

請求項 4 に記載の音高変換装置において、

重み関数格納ユニットが、 $f_a < f < f_b$ の区間は、 $W(f)$ の値が周波数 f の対数値に対して反比例する値となる重み関数 $W(f)$ を格納していることを特徴とする符号化音声データの音高変換装置。

【請求項 6】

請求項 2 または 3 に記載の音高変換装置において、

変換対象データ入力ユニットが、周波数 f をノートナンバー n によって示す符号を含む

符号化音声データDを入力し、

重み関数格納ユニットが、ノートナンバー n について定義された重み関数 $W(n)$ を格納し、

オフセット値入力ユニットが、音高を高める場合は正、低める場合は負の値をとるノートナンバーの差をオフセット値 α として入力し、

音高変換処理ユニットが、変換対象となる符号化音声データDに含まれる個々の符号について、当該符号が示すノートナンバー n を用いた「 $n' = n + \alpha \cdot W(n)$ 」なる演算式により新たなノートナンバー n' を求め、当該符号を、それが示すノートナンバー n を n' に変更した新たな符号に置き換える処理を行うことを特徴とする符号化音声データの音高変換装置。

10

【請求項7】

請求項6に記載の音高変換装置において、

重み関数格納ユニットが、周波数 f 軸上の第1設定値 f_a に対応するノートナンバー n_a および第2設定値 f_b に対応するノートナンバー n_b について、 $n_a < n < n_b$ の区間は、 $W(n)$ の値がノートナンバー n に反比例する値となる重み関数 $W(n)$ を格納していることを特徴とする符号化音声データの音高変換装置。

【請求項8】

請求項1～7のいずれかに記載の音高変換装置において、

変換対象データ入力ユニットが、音の持続時間が時間軸上で同一期間を占め、互いに異なる周波数を示す複数の符号を含む符号化音声データDを入力し、

20

音高変換処理ユニットが、前記複数の符号のそれぞれについて新たな符号への置換処理を行うことを特徴とする符号化音声データの音高変換装置。

【請求項9】

請求項8に記載の音高変換装置において、

変換対象データ入力ユニットが、周波数をノートナンバーによって示す符号を含む符号化音声データDを入力し、

音高変換処理ユニットが、音の持続時間が時間軸上で同一期間を占める複数の符号についてそれぞれ新たな符号への置換処理を行う際に、同一のノートナンバーを示す新たな符号が複数 m 個生じた場合には、当該複数 m 個の符号のうち1つのみを残し、その余の $(m-1)$ 個を削除する重複回避処理を行うことを特徴とする符号化音声データの音高変換装置。

30

【請求項10】

請求項9に記載の音高変換装置において、

変換対象データ入力ユニットが、音の強度の情報をもった符号を含む符号化音声データDを入力し、

音高変換処理ユニットが、重複回避処理を行う際に、1つのみ残された符号についての強度を、削除された符号についての強度に応じて修正することを特徴とする符号化音声データの音高変換装置。

【請求項11】

請求項1～10のいずれかに記載の音高変換装置において、

40

変換対象データ入力ユニットが、符号化音声データDとしてMIDI規格のデータを入力し、

音高変換処理ユニットが、変更後の符号化音声データ D^* としてMIDI規格のデータを出力することを特徴とする符号化音声データの音高変換装置。

【請求項12】

請求項11に記載の音高変換装置において、

音高変換処理ユニットが、変更後の符号化音声データ D^* を、五線譜上に音符を配置した楽譜のデータとして出力することを特徴とする符号化音声データの音高変換装置。

【請求項13】

請求項1～12のいずれかに記載の音高変換装置を含む音声変換装置であって、

50

人間の声を含む音声信号 S をアナログ信号もしくはデジタル信号として入力する音声信号入力ユニットと、

前記音声信号 S を、特定周波数の音が特定時間だけ持続することを示す符号を時間軸上に並べることによって構成された符号化音声データ D に変換する符号化ユニットと、

を有する音声符号化装置を更に備え、

前記符号化ユニットによって変換された符号化音声データ D の音高を、前記音高変換装置によって変更し、変更後の符号化音声データ D * を出力することを特徴とする音声変換装置。

【請求項 1 4】

請求項 1 ～ 1 2 のいずれかに記載の音高変換装置を含む音声合成装置であって、

所定の言語による単語を文字列として表現したテキストデータを入力するテキストデータ入力ユニットと、

前記所定の言語による単語を構成する個々の音節にそれぞれ対応する符号群（特定周波数の音が特定時間だけ持続することを示す符号の集合体）を格納した符号データベースユニットと、

前記符号データベースユニットを参照して、前記テキストデータの読みを構成する個々の音節にそれぞれ対応する符号群を抽出し、これらを時間軸上に並べることによって、前記テキストデータの読みに対応する人間の声を表現した符号化音声データ D を合成し、これを前記音高変換装置に与える符号合成ユニットと、

を有するテキスト符号化装置を更に備え、

前記符号合成ユニットによって合成された符号化音声データ D の音高を、前記音高変換装置によって変更し、変更後の符号化音声データ D * を出力することを特徴とする音声合成装置。

【請求項 1 5】

請求項 1 4 に記載の音声合成装置であって、

符号データベースユニットが、子音を構成する子音音素と母音を構成する母音音素とについて、それぞれ対応する符号群を格納しており、

符号合成ユニットが、テキストデータの読みを構成する個々の音節を子音音素と母音音素とに分解し、個々の音素ごとにそれぞれ対応する符号群を抽出し、これらを時間軸上に並べることによって符号化音声データ D を合成することを特徴とする音声合成装置。

【請求項 1 6】

請求項 1 4 または 1 5 に記載の音声合成装置であって、

テキストデータ入力ユニットが、音節ごとのオフセット値 α を含むテキストデータを入力し、

符号合成ユニットが、合成した符号化音声データ D とともに、個々の音節ごとの前記オフセット値 α を音高変換装置に与え、

前記音高変換装置が、前記符号合成ユニットから与えられたオフセット値 α を用いて、個々の音節ごとに音高の変更処理を行うことを特徴とする音声合成装置。

【請求項 1 7】

請求項 1 4 または 1 5 に記載の音声合成装置であって、

テキスト符号化装置が、

個々の単語について、当該単語を構成する各音節に与えるオフセット値 α を格納した音高辞書ユニットを更に備え、

符号合成ユニットが、合成した符号化音声データ D とともに、前記音高辞書ユニットを参照することにより得られる個々の音節ごとの前記オフセット値 α を音高変換装置に与え、

前記音高変換装置が、前記符号合成ユニットから与えられたオフセット値 α を用いて、個々の音節ごとに音高の変更処理を行うことを特徴とする音声合成装置。

【請求項 1 8】

請求項 1 3 ～ 1 7 のいずれかに記載の音声変換装置もしくは音声合成装置であって、

所定の楽器による様々な周波数の演奏音響波形をデジタルデータとして格納した音源ユニットと、

符号化音声データ D^* を構成する個々の符号を、前記音源ユニットに格納されている対応する演奏音響波形に置き換えることにより音声信号の復号化を行う復号化ユニットと、
復号化された音声信号に基づいて音波を生成する発音ユニットと、

を有する音声発生装置を更に備えることを特徴とする音声変換装置もしくは音声合成装置。

【請求項 19】

請求項 1 ～ 18 のいずれかに記載の音高変換装置、音声変換装置もしくは音声合成装置としてコンピュータを機能させるためのプログラム。

10

【請求項 20】

特定周波数の音が特定時間だけ持続することを示す符号を時間軸上に並べることによって構成され人間の声を表現する符号化音声データについて、その抑揚を変換する符号化音声データの抑揚変換方法であって、

コンピュータが、変換対象となる符号化音声データ D を入力する変換対象データ入力段階と、

コンピュータが、音高に関するオフセット値 α を入力するオフセット値入力段階と、

コンピュータが、前記変換対象となる符号化音声データ D に対して、前記オフセット値 α に基づく音高の変更処理を行い、変更後の符号化音声データ D^* を出力する音高変換処理段階と、

20

を有し、

前記音高変換処理段階において、周波数 f 軸上の所定区間において $W(f)$ が周波数 f の増加に従って単調減少する所定の重み関数 $W(f)$ を利用して、前記変換対象となる符号化音声データ D に含まれている個々の符号について、当該符号が示す周波数 f を $\alpha \cdot W(f)$ に応じた値だけ増減することにより新たな周波数 f' を求め、当該符号を、それが示す周波数 f を f' に変更した新たな符号に置き換える処理を行うことを特徴とする符号化音声データの抑揚変換方法。

【請求項 21】

人間の声を表現した M I D I データを、抑揚の異なる別な M I D I データに変換する M I D I データの抑揚変換方法であって、

30

コンピュータが、変換対象となる M I D I データ D を入力する変換対象データ入力段階と、

コンピュータが、音高に関するオフセット値 α を入力するオフセット値入力段階と、

コンピュータが、前記変換対象となる M I D I データ D に対して、前記オフセット値 α に基づく音高の変更処理を行い、変更後の M I D I データ D^* を出力する音高変換処理段階と、

を有し、

前記変換対象データ入力段階では、互いに異なるノートナンバーをもち、時間軸上の同一位置を占める複数の M I D I 符号を含む M I D I データ D の入力を行い、

前記音高変換処理段階では、ノートナンバー軸上の第 1 設定値 n_a および第 2 設定値 n_b ($n_b > n_a$) について、 $n \leq n_a$ の区間は、 $W(n) = 1$ 、 $n_a < n < n_b$ の区間は、 $1 > W(n) > 0$ (但し、 $W(n)$ は n の増加に従って単調減少)、 $n \geq n_b$ の区間は、 $W(n) = 0$ となる所定の重み関数 $W(n)$ を利用して、前記変換対象となる M I D I データ D に含まれている個々の M I D I 符号について、当該 M I D I 符号が示すノートナンバー n に対して $\alpha \cdot W(n)$ に応じた値だけ加減算を行うことにより新たなノートナンバー n' を求め、当該 M I D I 符号を、それが示すノートナンバー n を n' に変更した新たな M I D I 符号に置き換える処理を行うことにより、変更後の M I D I データ D^* を生成することを特徴とする M I D I データの抑揚変換方法。

40

【請求項 22】

請求項 21 に記載の M I D I データの抑揚変換方法において、

50

音高変換処理段階で、新たなMIDI符号に置き換える処理を行う際に、時間軸上の同一位置を占め、同一のノートナンバーをもつ新たなMIDI符号が複数 m 個生じた場合には、当該複数 m 個のMIDI符号のうち1つのみを残し、その余の $(m-1)$ 個を削除する重複回避処理を行うことを特徴とするMIDIデータの抑揚変換方法。

【発明の詳細な説明】

【技術分野】

【0001】

本発明は、符号化音声データの音高変換装置およびその応用技術に関し、特に、人間の声を表現したMIDI規格等の符号化音声データを、異なる音高をもった別な音声データに変換する技術に関する。

【背景技術】

【0002】

音響信号を符号化する技術としては、PCM(Pulse Code Modulation)の手法が古くから知られており、現在、コンピュータをはじめとする多くのデジタル機器で利用されているデジタル音響データは、このPCMの手法を用いてデジタル化したデータである。

【0003】

一方、最近では、パーソナルコンピュータの普及とともに、MIDI規格(Musical Instrument Digital Interface)による符号化も一般化してきている。もともとMIDI規格は、電子楽器による楽器音を符号化しようという発想から生まれた規格であり、PCMとは全く異なる固有の特徴を有している。すなわち、MIDI規格による符号データ(以下、MIDIデータという)は、基本的には、楽器のどの鍵盤キーを、どの程度の強さで弾いたか、という楽器演奏の操作を記述したデータであり、このMIDIデータ自身には、実際の音の波形は含まれていない。そのため、実際の音を再生する場合には、楽器音の波形を記憶したMIDI音源が別途必要になる。しかしながら、上述したPCMの手法で音を記録する場合に比べて、情報量が極めて少なくすむという特徴を有し、その符号化効率の高さが注目を集めている。

【0004】

最近では、楽器音に限らず、人間の声などを含めた様々な音響信号をMIDIデータによって符号化しようとする試みもなされている。たとえば、下記の特許文献1および2には、MIDIデータを利用することが可能な新規な符号化方法が提案されている。これらの方法では、音響信号の時間軸に沿って複数の単位区間を設定し、各単位区間ごとにフーリエ変換を行ってスペクトルを求め、このスペクトルに応じたMIDIデータを作成するという手順が実行される。また、下記の特許文献3には、人間の声や歌声を含む、いわゆるヴォーカル音響信号について、MIDIデータを作成する効率的な手法が提案されている。更に、下記の特許文献4には、日本語のカナ文字を構成する各音節を符号コードで表現し、人間の声を基にしてMIDIデータを作成する技術が提案されている。

【先行技術文献】

【特許文献】

【0005】

【特許文献1】特開平11-95753号公報

【特許文献2】特開2000-99009号公報

【特許文献3】特開2000-99093号公報

【特許文献4】特願2009-244698号明細書

【発明の概要】

【発明が解決しようとする課題】

【0006】

人間の話し声には抑揚があり、同じ音節からなる言葉でも、抑揚を変化させてしゃべると、それぞれ違った印象をもった言葉に聞こえてくる。また、人間の歌声、いわゆるヴォーカル音にも、当然ながら、音の高低が存在する。したがって、人間の話し声や歌声を表現した符号化音声データについて、その音高(ピッチ)を自由に改変することができれば

10

20

30

40

50

、その応用範囲は多岐に広がることになる。たとえば、人間の話し声に対して、音節単位で音高を上下させることができれば、言葉の抑揚を変化させたり、所望の曲に合わせた歌声を作成したりすることができる。

【0007】

MIDIデータなどの符号化データに対する音高シフトは、音楽でいう、いわゆる移調演奏に相当するものであり、通常、すべての音符のノートナンバーに対して、一律に所定のオフセット値を加減算することによって行われる。たとえば、一般的な音楽MIDIデータについて、1オクターブだけ上げる音高シフトを行うのであれば、すべての音符のノートナンバー n に対して、1オクターブの音程に相当するオフセット値「12」を加える加算処理を行えばよい。

10

【0008】

しかしながら、本願発明者が行った実験によると、人間の声を表現したMIDIデータについて、同様の方法で音高シフトを行うと、音声の明瞭性は維持されるものの、声質が破壊され、違和感のある奇怪な声になってしまうことが判明した。たとえば、人間の声を表現したMIDIデータについて、1オクターブだけ上げる音高シフトを行うために、すべての音符のノートナンバー n に対して、一律にオフセット値「12」を加える加算処理を実行すると、個々の単語の判別は可能であるものの、いわゆるボイスチェンジャーを通したような違和感のある不自然な声になってしまう。

【0009】

本願発明者は、当初、人間の声を表現したMIDI規格等の符号化音声データに対して、所望の音高シフト処理を施し、違和感のない自然な音声再生が可能な符号化音声データを生成するためには、人間の声に含まれる音声成分に対する高度な信号解析技術を導入する必要があると考えていた。しかしながら、そのような高度な音声信号解析技術を導入した処理を実行するためには、高性能なハードウェアを用いた複雑な処理が必要となり、そのような処理機能をもった音高変換装置は極めてコストの高い装置にならざるを得ない。

20

【0010】

そこで本発明は、人間の声を表現した符号化音声データに対して、できるだけ簡便な方法により所望の音高シフト処理を施し、違和感のない自然な音声再生が可能な符号化音声データを取得することができる符号化音声データの音高変換装置および符号化音声データの抑揚変換方法を提供することを目的とする。

30

【0011】

また、本発明は、上記音高変換装置を利用して、入力した人間の音声の音高を変更し、違和感のない自然な音声再生が可能な符号化音声データを取得することができる音声変換装置を提供することを目的とし、更に、所定の言語による文章を文字列として表現したテキストデータに基づいて、自然な抑揚をもった音声再生が可能な符号化音声データを取得することができる音声合成装置を提供することを目的とする。

【課題を解決するための手段】

【0012】

(1) 本発明の第1の態様は、特定周波数の音が特定時間だけ持続することを示す符号を時間軸上に並べることによって構成され人間の声を表現する符号化音声データを、異なる音高をもった別な音声データに変換する符号化音声データの音高変換装置において、変換対象となる符号化音声データDを入力する変換対象データ入力ユニットと、音高に関するオフセット値 α を入力するオフセット値入力ユニットと、周波数 f について定義された所定の重み関数 $W(f)$ を格納した重み関数格納ユニットと、

40

変換対象となる符号化音声データDに対して、オフセット値 α に基づく音高の変更処理を行い、変更後の符号化音声データ D^* を出力する音高変換処理ユニットと、を設け、

重み関数格納ユニットが、重み関数 $W(f)$ として、周波数 f 軸上の所定区間において $W(f)$ が周波数 f の増加に従って単調減少する関数を格納しており、

50

音高変換処理ユニットが、重み関数 $W(f)$ を用いて、変換対象となる符号化音声データ D に含まれている個々の符号について、当該符号が示す周波数 f を $\alpha \cdot W(f)$ に応じた値だけ増減することにより新たな周波数 f' を求め、当該符号を、それが示す周波数 f を f' に変更した新たな符号に置き換える処理を行うようにしたものである。

【0013】

(2) 本発明の第2の態様は、上述の第1の態様に係る符号化音声データの音高変換装置において、

重み関数格納ユニットが、重み関数 $W(f)$ として、周波数 f 軸上の第1設定値 f_a および第2設定値 f_b ($f_b > f_a$) について、 $f \leq f_a$ の区間は、 $W(f) = 1$ 、 $f_a < f < f_b$ の区間は、 $1 > W(f) > 0$ (但し、 $W(f)$ は f の増加に従って単調減少)、 $f \geq f_b$ の区間は、 $W(f) = 0$ となる関数を格納しているようにしたものである。

10

【0014】

(3) 本発明の第3の態様は、上述の第2の態様に係る符号化音声データの音高変換装置において、

重み関数格納ユニットが、100Hz～200Hzの範囲内の第1設定値 f_a と3kHz～6kHzの範囲内の第2設定値 f_b とを用いた男性用重み関数 $W\text{-male}(f)$ と、200Hz～400Hzの範囲内の第1設定値 f_a と4kHz～8kHzの範囲内の第2設定値 f_b とを用いた女性用重み関数 $W\text{-female}(f)$ と、を格納しており、

音高変換処理ユニットが、変換対象となる符号化音声データ D によって表現される声が男性の声か女性の声かを示す指示に基づいて、男性の声の場合には男性用重み関数 $W\text{-male}(f)$ を用いた変更処理を行い、女性の声の場合には女性用重み関数 $W\text{-female}(f)$ を用いた変更処理を行うようにしたものである。

20

【0015】

(4) 本発明の第4の態様は、上述の第2または第3の態様に係る符号化音声データの音高変換装置において、

オフセット値入力ユニットが、音高を高める場合は正、低める場合は負のオフセット値 α を入力し、

音高変換処理ユニットが、所定の係数 k ($k > 1$) を用いた式 $f' = f \cdot k^{\alpha \cdot W(f)}$ により新たな周波数 f' を求めるようにしたものである。

【0016】

30

(5) 本発明の第5の態様は、上述の第4の態様に係る符号化音声データの音高変換装置において、

重み関数格納ユニットが、 $f_a < f < f_b$ の区間は、 $W(f)$ の値が周波数 f の対数値に対して反比例する値となる重み関数 $W(f)$ を格納しているようにしたものである。

【0017】

(6) 本発明の第6の態様は、上述の第2または第3の態様に係る符号化音声データの音高変換装置において、

変換対象データ入力ユニットが、周波数 f をノートナンバー n によって示す符号を含む符号化音声データ D を入力し、

重み関数格納ユニットが、ノートナンバー n について定義された重み関数 $W(n)$ を格納し、

40

オフセット値入力ユニットが、音高を高める場合は正、低める場合は負の値をとるノートナンバーの差をオフセット値 α として入力し、

音高変換処理ユニットが、変換対象となる符号化音声データ D に含まれる個々の符号について、当該符号が示すノートナンバー n を用いた「 $n' = n + \alpha \cdot W(n)$ 」なる演算式により新たなノートナンバー n' を求め、当該符号を、それが示すノートナンバー n を n' に変更した新たな符号に置き換える処理を行うようにしたものである。

【0018】

(7) 本発明の第7の態様は、上述の第6の態様に係る符号化音声データの音高変換装置において、

50

重み関数格納ユニットが、周波数 f 軸上の第 1 設定値 f_a に対応するノートナンバー n_a および第 2 設定値 f_b に対応するノートナンバー n_b について、 $n_a < n < n_b$ の区間は、 $W(n)$ の値がノートナンバー n に反比例する値となる重み関数 $W(n)$ を格納しているようにしたものである。

【0019】

(8) 本発明の第 8 の態様は、上述の第 1 ～第 7 の態様に係る符号化音声データの音高変換装置において、

変換対象データ入力ユニットが、音の持続時間が時間軸上で同一期間を占め、互いに異なる周波数を示す複数の符号を含む符号化音声データ D を入力し、

音高変換処理ユニットが、複数の符号のそれぞれについて新たな符号への置換処理を行うようにしたものである。

10

【0020】

(9) 本発明の第 9 の態様は、上述の第 8 の態様に係る符号化音声データの音高変換装置において、

変換対象データ入力ユニットが、周波数をノートナンバーによって示す符号を含む符号化音声データ D を入力し、

音高変換処理ユニットが、音の持続時間が時間軸上で同一期間を占める複数の符号についてそれぞれ新たな符号への置換処理を行う際に、同一のノートナンバーを示す新たな符号が複数 m 個生じた場合には、当該複数 m 個の符号のうち 1 つのみを残し、その余の $(m - 1)$ 個を削除する重複回避処理を行うようにしたものである。

20

【0021】

(10) 本発明の第 10 の態様は、上述の第 9 の態様に係る符号化音声データの音高変換装置において、

変換対象データ入力ユニットが、音の強度の情報をもった符号を含む符号化音声データ D を入力し、

音高変換処理ユニットが、重複回避処理を行う際に、1 つのみ残された符号についての強度を、削除された符号についての強度に応じて修正するようにしたものである。

【0022】

(11) 本発明の第 11 の態様は、上述の第 1 ～第 10 の態様に係る符号化音声データの音高変換装置において、

30

変換対象データ入力ユニットが、符号化音声データ D として M I D I 規格のデータを入力し、

音高変換処理ユニットが、変更後の符号化音声データ D^* として M I D I 規格のデータを出力するようにしたものである。

【0023】

(12) 本発明の第 12 の態様は、上述の第 11 の態様に係る符号化音声データの音高変換装置において、

音高変換処理ユニットが、変更後の符号化音声データ D^* を、五線譜上に音符を配置した楽譜のデータとして出力するようにしたものである。

【0024】

40

(13) 本発明の第 13 の態様は、上述の第 1 ～第 12 の態様に係る符号化音声データの音高変換装置に、

人間の声を含む音声信号 S をアナログ信号もしくはデジタル信号として入力する音声信号入力ユニットと、

音声信号 S を、特定周波数の音が特定時間だけ持続することを示す符号を時間軸上に並べることによって構成された符号化音声データ D に変換する符号化ユニットと、

を有する音声符号化装置を更に付加し、

符号化ユニットによって変換された符号化音声データ D の音高を、音高変換装置によって変更し、変更後の符号化音声データ D^* を出力する機能をもった音声変換装置を構成するようにしたものである。

50

【0025】

(14) 本発明の第14の態様は、上述の第1～第12の態様に係る符号化音声データの音高変換装置に、

所定の言語による単語を文字列として表現したテキストデータを入力するテキストデータ入力ユニットと、

所定の言語による単語を構成する個々の音節にそれぞれ対応する符号群（特定周波数の音が特定時間だけ持続することを示す符号の集合体）を格納した符号データベースユニットと、

符号データベースユニットを参照して、テキストデータの読みを構成する個々の音節にそれぞれ対応する符号群を抽出し、これらを時間軸上に並べることによって、テキストデータの読みに対応する人間の声を表現した符号化音声データDを合成し、これを音高変換装置に与える符号合成ユニットと、

を有するテキスト符号化装置を更に付加し、

符号合成ユニットによって合成された符号化音声データDの音高を、音高変換装置によって変更し、変更後の符号化音声データD*を出力する機能をもった音声合成装置を構成するようにしたものである。

【0026】

(15) 本発明の第15の態様は、上述の第14の態様に係る音声合成装置において、

符号データベースユニットが、子音を構成する子音音素と母音を構成する母音音素とについて、それぞれ対応する符号群を格納しており、

符号合成ユニットが、テキストデータの読みを構成する個々の音節を子音音素と母音音素とに分解し、個々の音素ごとにそれぞれ対応する符号群を抽出し、これらを時間軸上に並べることによって符号化音声データDを合成するようにしたものである。

【0027】

(16) 本発明の第16の態様は、上述の第14または第15の態様に係る音声合成装置において、

テキストデータ入力ユニットが、音節ごとのオフセット値 α を含むテキストデータを入力し、

符号合成ユニットが、合成した符号化音声データDとともに、個々の音節ごとのオフセット値 α を音高変換装置に与え、

音高変換装置が、符号合成ユニットから与えられたオフセット値 α を用いて、個々の音節ごとに音高の変更処理を行うようにしたものである。

【0028】

(17) 本発明の第17の態様は、上述の第14または第15の態様に係る音声合成装置において、

テキスト符号化装置が、

個々の単語について、当該単語を構成する各音節に与えるオフセット値 α を格納した音高辞書ユニットを更に備え、

符号合成ユニットが、合成した符号化音声データDとともに、音高辞書ユニットを参照することにより得られる個々の音節ごとのオフセット値 α を音高変換装置に与え、

音高変換装置が、符号合成ユニットから与えられたオフセット値 α を用いて、個々の音節ごとに音高の変更処理を行うようにしたものである。

【0029】

(18) 本発明の第18の態様は、上述の第13～第17の態様に係る音声変換装置もしくは音声合成装置において、

所定の楽器による様々な周波数の演奏音響波形をデジタルデータとして格納した音源ユニットと、

符号化音声データD*を構成する個々の符号を、音源ユニットに格納されている対応する演奏音響波形に置き換えることにより音声信号の復号化を行う復号化ユニットと、

復号化された音声信号に基づいて音波を生成する発音ユニットと、

を有する音声発生装置を更に設けるようにしたものである。

【0030】

(19) 本発明の第19の態様は、上述の第1～第18の態様に係る音高変換装置、音声変換装置もしくは音声合成装置を、コンピュータにプログラムを組み込むことにより構成したものである。

【0031】

(20) 本発明の第20の態様は、特定周波数の音が特定時間だけ持続することを示す符号を時間軸上に並べることによって構成され人間の声を表現する符号化音声データについて、その抑揚を変換する符号化音声データの抑揚変換方法において、

コンピュータが、変換対象となる符号化音声データDを入力する変換対象データ入力段階と、

コンピュータが、音高に関するオフセット値 α を入力するオフセット値入力段階と、

コンピュータが、変換対象となる符号化音声データDに対して、オフセット値 α に基づく音高の変更処理を行い、変更後の符号化音声データD*を出力する音高変換処理段階と

、

を行い、

音高変換処理段階において、周波数 f 軸上の所定区間において $W(f)$ が周波数 f の増加に従って単調減少する所定の重み関数 $W(f)$ を利用して、変換対象となる符号化音声データDに含まれている個々の符号について、当該符号が示す周波数 f を $\alpha \cdot W(f)$ に応じた値だけ増減することにより新たな周波数 f' を求め、当該符号を、それが示す周波数 f を f' に変更した新たな符号に置き換える処理を行うようにしたものである。

【0032】

(21) 本発明の第21の態様は、人間の声を表現したMIDIデータを、抑揚の異なる別なMIDIデータに変換するMIDIデータの抑揚変換方法において、

コンピュータが、変換対象となるMIDIデータDを入力する変換対象データ入力段階と、

コンピュータが、音高に関するオフセット値 α を入力するオフセット値入力段階と、

コンピュータが、変換対象となるMIDIデータDに対して、オフセット値 α に基づく音高の変更処理を行い、変更後のMIDIデータD*を出力する音高変換処理段階と、

を行い、

変換対象データ入力段階では、互いに異なるノートナンバーをもち、時間軸上の同一位置を占める複数のMIDI符号を含むMIDIデータDの入力を行い、

音高変換処理段階では、ノートナンバー軸上の第1設定値 n_a および第2設定値 n_b ($n_b > n_a$) について、 $n \leq n_a$ の区間は、 $W(n) = 1$ 、 $n_a < n < n_b$ の区間は、 $1 > W(n) > 0$ (但し、 $W(n)$ は n の増加に従って単調減少)、 $n \geq n_b$ の区間は、 $W(n) = 0$ となる所定の重み関数 $W(n)$ を利用して、変換対象となるMIDIデータDに含まれている個々のMIDI符号について、当該MIDI符号が示すノートナンバー n に対して $\alpha \cdot W(n)$ に応じた値だけ加減算を行うことにより新たなノートナンバー n' を求め、当該MIDI符号を、それが示すノートナンバー n を n' に変更した新たなMIDI符号に置き換える処理を行うことにより、変更後のMIDIデータD*を生成するよう

【0033】

(22) 本発明の第22の態様は、上述の第21の態様に係るMIDIデータの抑揚変換方法において、

音高変換処理段階で、新たなMIDI符号に置き換える処理を行う際に、時間軸上の同一位置を占め、同一のノートナンバーをもつ新たなMIDI符号が複数 m 個生じた場合には、当該複数 m 個のMIDI符号のうち1つのみを残し、その余の $(m-1)$ 個を削除する重複回避処理を行うようにしたものである。

【発明の効果】

【0034】

10

20

30

40

50

本発明に係る符号化音声データの音高変換装置および符号化音声データの抑揚変換方法によれば、変換対象となる符号化音声データDに対して、オフセット値 α に基づく音高の変更処理を行って符号化音声データD*を得る際に、周波数 f の増加に従って単調減少する重み関数 $W(f)$ を用いて、符号が示す周波数 f を $\alpha \cdot W(f)$ に応じた値だけ増減するようにしたため、より高音の符号ほど音高シフト量が減少することになる。このため、人間の発生する音声に含まれる基本周波数成分の音高シフト量に比べて、当該基本周波数成分の整数倍の周波数をもつフォルマント成分についての音高シフト量が小さくなり、「フォルマント成分の絶対周波数値はほぼ一定」という人間の声音に固有の特徴を維持したまま、音高を変更することができる。かくして、比較的簡便な方法で、人間の声を表現したMIDI規格等の符号化音声データを、異なる音高、抑揚をもった別な音声データに変換することができるようになり、しかも違和感のない自然な音声再生が可能な符号化音声データが得られる。

10

【0035】

なお、本発明で取り扱う符号化音声データは、たとえばMIDIデータのように、特定周波数の音が特定時間だけ持続することを示すいくつかの符号を時間軸上に並べることによって構成されるデータであり、いわば人間の声を音符で表現したデータというべきものである。したがって、そのような符号化音声データを再生したとしても、当然ながら、元の人間の声に忠実な音が生成されるわけではなく、人間の声に対して、かなりかけ離れた音が得られることになる。したがって、人間の声に忠実なリアルな音か否かという観点では、当然ながら、本発明は「違和感のない自然な音声再生」を可能にする技術ではない。本発明の効果にいう「違和感のない自然な音声再生」とは、「人間の声をもつフォルマント成分に起因して生じる固有の特徴が損なわれることのない再生」を意味するものであり、いわゆるボイスチェンジャーを通した場合に得られるような違和感や不自然さが排除されることを意味するものである。

20

【0036】

本発明では、特に、平均的な音声についての基本周波数成分(F_0 : 0次フォルマント成分)が含まれると予想される周波数領域の上限に第1設定値 f_a を定義し、聴取可能な最高次フォルマント $F^{\max}$ 成分が含まれると予想される周波数領域の上限に第2設定値 f_b を定義し、 $f \leq f_a$ の区間は、 $W(f) = 1$ 、 $f_a < f < f_b$ の区間は、 $1 > W(f) > 0$ (但し、 $W(f)$ は f の増加に従って単調減少)、 $f \geq f_b$ の区間は、 $W(f) = 0$ となる関数を重み関数 $W(f)$ として設定すれば、非常に単純な演算によって音高の変更処理を行うことができる。特に、MIDIデータのように、周波数 f をその対数値に対応するノートナンバー n で表す符号データを用いる場合、重み関数として n の関数 $W(n)$ を用いるようにすれば、音高の変更処理は、ノートナンバー n に対して、 $\alpha \cdot W(n)$ を加算もしくは減算するだけの単純な処理になる。上記周波数 $f_a < f < f_b$ の区間について、ノートナンバー n に反比例する値をとる関数値 $W(n)$ を用いるようにすれば、演算はより単純になる。

30

【0037】

なお、MIDI規格等の符号化音声データによって人間の声を表すと、音の持続時間が時間軸上で同一期間を占め、互いに異なる周波数を示す複数の符号(いわゆる、和音を構成する複数の符号)を含むデータが生じることになるが、このような和音を構成する複数の符号に対して音高の変更処理を行うと、複数の符号が変更後に同じ音高を占める場合がありえる。このように同じ音高の符号が重複した場合には、1つのみを残して他を削除するようにすれば、重複を許さないMIDI規格のような符号化データに対しても問題なく適用可能になる。また、削除された符号についての強度に応じて強度修正を行えば、音高の変更処理後も適切な強度バランスをもった符号化データが得られる。

40

【0038】

また、上記音高変換装置を利用すれば、入力した人間の音声の音高を自由に変更して、違和感のない自然な音声再生が可能な符号化音声データを生成する音声変換装置を実現することができる。更に、所定の言語による文章を文字列として表現したテキストデータに

50

基づいて、自然な抑揚をもった音声再生が可能な符号化音声データを得ることができる音声合成装置を実現することもできる。

【図面の簡単な説明】

【0039】

【図1】フーリエ変換を利用した音響信号の符号化方法の基本原理を示す図である。

【図2】図1(c)に示す強度グラフに基いて作成された符号コードを示す図である。

【図3】時間軸上に部分的に重複するように単位区間設定を行うことにより作成された符号コードを示す図である。

【図4】和音に対して音高を1オクターブだけ上げる一般的な処理を示す楽譜である。

【図5】人間の声に含まれるフォルマント成分を例示する強度グラフ（人間の音声スペクトルのピーク位置を示すグラフ）である。

【図6】図(a)は、標準的な音高で発声を行った人間の音声スペクトルの強度グラフ、図(b)は、同じ人が音高を高めて発声を行った場合の強度グラフ、図(c)は、同じ人が音高を低めて発声を行った場合の強度グラフである。

【図7】図(a)は、本発明で用いる、周波数 f についての重み関数 $W(f)$ の一例を示すグラフ（横軸は、周波数 f についての線形スケール）、図(b)は、本発明で用いる、ノートナンバー n についての重み関数 $W(n)$ の一例を示すグラフ（横軸は、周波数 f についての対数スケール）である。

【図8】図(a)は、変換前の符号化データに含まれる和音の一例を示す楽譜、図(b)は、従来の方法で音高を1オクターブ上げる処理を行った結果を示す楽譜、図(c)は、本発明の方法で音高を1オクターブ上げる処理を行った結果を示す楽譜である。

【図9】図(a)は、変換前の符号化データに含まれる和音の一例を示す楽譜、図(b)は、従来の方法で音高を1オクターブ下げる処理を行った結果を示す楽譜、図(c)は、本発明の方法で音高を1オクターブ下げる処理を行った結果を示す楽譜である。

【図10】図(a)は、変換前の符号化データに含まれる和音の一例を示す楽譜、図(b)は、従来の方法で音高を1オクターブ上げる処理を行った結果を示す楽譜、図(c)は、本発明の方法で音高を1オクターブ上げる処理を行った結果（3つの音符が重なっている状態）を示す楽譜である。

【図11】音高変更後に複数の符号が同じ音高を占める場合に、1つのみを残して他を削除する処理を示す楽譜である。

【図12】本発明の基本的実施形態に係る符号化音声データの音高変換装置の構成を示すブロック図である。

【図13】図12に示す音高変換装置を利用した音声変換装置の構成を示すブロック図である。

【図14】図12に示す音高変換装置を利用した音声合成装置の構成を示すブロック図である。

【図15】図14に示す音声合成装置の符号合成ユニット330内に用意されている変換テーブルを示す図である。

【図16】図15に示す変換テーブルに基づいて、テキストデータを構成する個々の音節を子音音素と母音音素とに分解した例を示す図である。

【図17】図14に示す音声合成装置の符号データベースユニット310内に用意されている男性用のMIDI符号データベースの一例を示す表である。

【図18】図14に示す音声合成装置の符号データベースユニット310内に用意されている女性用のMIDI符号データベースの一例を示す表である。

【図19】図17に示す男性用のMIDI符号データベースを用いて、各母音音素と各子音音素とに対応する符号を、音符として表示した例を示す楽譜である。

【図20】図18に示す女性用のMIDI符号データベースを用いて、各母音音素と各子音音素とに対応する符号を、音符として表示した例を示す楽譜である。

【図21】図14に示す音声合成装置の第1の変形例を示すブロック図である。

【図22】図21に示す音声合成装置の符号合成ユニット340の処理機能を説明する図

10

20

30

40

50

である。

【図 2 3】図 1 4 に示す音声合成装置の第 2 の変形例を示すブロック図である。

【図 2 4】図 2 3 に示す音声合成装置の音高辞書ユニット 3 6 0 の内容を例示する表である。

【図 2 5】本発明に係る音高変換装置から出力された符号化音声データ D * に基づいて、音声を発声させる音声発声装置の基本構成を示すブロック図である。

【発明を実施するための形態】

【0 0 4 0】

以下、本発明を図示する実施形態に基づいて説明する。

【0 0 4 1】

<<< § 1. 音響信号の符号化方法の基本原則 >>>

はじめに、前掲の特許文献 1 ～ 3 に開示されているフーリエ変換を利用した音響信号の符号化方法の基本原則を簡単に説明しておく。この符号化方法を利用すれば、たとえば、人間の音声をアナログ音響信号として取り込み、M I D I データなどのデジタル符号データに変換することができる。

【0 0 4 2】

いま、図 1 (a) に示すように、時系列の強度信号としてアナログ音響信号が与えられたものとしよう。図示の例では、横軸に時間 t 、縦軸に振幅（強度）をとってこの音響信号を示している。ここでは、まずこのアナログ音響信号を、デジタルの音響データとして取り込む処理を行う。これは、従来の一般的な P C M の手法を用い、所定のサンプリング周期でこのアナログ音響信号をサンプリングし、振幅を所定の量子化ビット数を用いてデジタルデータに変換する処理を行えばよい。ここでは、説明の便宜上、P C M の手法でデジタル化した音響データの波形も、図 1 (a) のアナログ音響信号と同一の波形で示すことにする。

【0 0 4 3】

続いて、この符号化対象となる音響信号の時間軸上に、複数の単位区間を設定する。図 1 (a) に示す例では、時間軸 t 上に等間隔に 6 つの時刻 $t_1 \sim t_6$ が定義され、これら各時刻を始点および終点とする 5 つの単位区間 $d_1 \sim d_5$ が設定されている。実際には、後述するように、各単位区間 $d_1 \sim d_5$ は部分的に重複するように設定するのが好ましい。

【0 0 4 4】

こうして単位区間が設定されたら、単位区間ごとの音響信号（ここでは、区間信号と呼ぶことにする）に対してそれぞれフーリエ変換を行い、スペクトルを作成する。このとき、ハニング窓（Hanning Window）などの重み関数で、切り出した区間信号にフィルタをかけてフーリエ変換を施す。一般にフーリエ変換は、切り出した区間前後に同様な信号が無限に存在することが想定されているため、重み関数を用いない場合、作成したスペクトルに高周波ノイズがのることが多い。ハニング窓関数など区間の両端の重みが 0 になるような重み関数を用いると、このような弊害をある程度抑制できる。ハニング窓関数 $H(k)$ は、単位区間長を L とすると、 $k = 1 \cdots L$ に対して、

$$H(k) = 0.5 - 0.5 \cdot \cos(2\pi k / L)$$

で与えられる関数である。

【0 0 4 5】

図 1 (b) には、単位区間 d_1 について作成されたスペクトルの一例が示されている。このスペクトルでは、横軸上に定義された周波数 f によって、単位区間 d_1 についての区間信号に含まれる周波数成分（ $0 \sim F$ ：ここで F はサンプリング周波数の $1/2$ ）が示されており、縦軸上に定義された複素強度 A によって、周波数成分ごとの複素強度が示されている。

【0 0 4 6】

次に、このスペクトルの周波数軸 f に対応させて、離散的に複数 X 個の符号コードを定義する。この例では、符号コードとして M I D I データで利用されるノートナンバー n を用いており、 $n = 0 \sim 127$ までの 128 個の符号コードを定義している。ノートナンバ

10

20

30

40

50

ー n は、音符の音階を示すパラメータであり、たとえば、ノートナンバー $n = 69$ は、ピアノの鍵盤中央の「ラ音（A 3 音）」を示しており、440 Hz の音に相当する。このように、128個のノートナンバーには、いずれも所定の周波数が対応づけられるので、スペクトルの周波数軸 f 上の所定位置に、それぞれ128個のノートナンバー n が離散的に定義されることになる。

【0047】

ここで、ノートナンバー n は、1オクターブ上がると、周波数が2倍になる対数尺度の音階を示すため、周波数軸 f に対して線形には対応しない。そこで、ここでは周波数軸 f を対数尺度で表し、この対数尺度軸上にノートナンバー n を定義した強度グラフを作成してみる。図1(c)は、このようにして作成された単位区間 $d1$ についての強度グラフを示す。この強度グラフの横軸は、図1(b)に示すスペクトルの横軸を対数尺度に変換したものであり、ノートナンバー $n = 0 \sim 127$ が等間隔にプロットされている。一方、この強度グラフの縦軸は、図1(b)に示すスペクトルの複素強度 A を実効強度 E に変換したものであり、各ノートナンバー n の位置における強度を示している。一般に、フーリエ変換によって得られる複素強度 A は、実数部 R と虚数部 I とによって表されるが、実効強度 E （エネルギー）は、 $E = \sqrt{(R^2 + I^2)}$ なる演算によって求めることができる。

【0048】

こうして求められた単位区間 $d1$ の強度グラフは、単位区間 $d1$ についての区間信号に含まれる振動成分について、ノートナンバー $n = 0 \sim 127$ に相当する各振動成分の割合を実効強度として示すグラフということができる。そこで、この強度グラフに示されている各実効強度に基いて、全 X 個（この例では $X = 128$ ）のノートナンバーの中から P 個のノートナンバーを選択し、この P 個のノートナンバー n を、単位区間 $d1$ を代表する代表符号コードとして抽出する。ここでは、説明の便宜上、 $P = 3$ として、全128個の候補の中から3個のノートナンバーを代表符号コードとして抽出する場合を示すことにする。たとえば、「候補の中から強度の大きい順に P 個の符号コードを抽出する」という基準に基いて抽出を行えば、図1(c)に示す例では、第1番目の代表符号コードとしてノートナンバー $n(d1, 1)$ が、第2番目の代表符号コードとしてノートナンバー $n(d1, 2)$ が、第3番目の代表符号コードとしてノートナンバー $n(d1, 3)$ が、それぞれ抽出されることになる。

【0049】

このようにして、 P 個の代表符号コードが抽出されたら、これらの代表符号コードとその実効強度によって、単位区間 $d1$ についての区間信号を表現することができる。たとえば、上述の例の場合、図1(c)に示す強度グラフにおいて、ノートナンバー $n(d1, 1)$ 、 $n(d1, 2)$ 、 $n(d1, 3)$ の実効強度がそれぞれ $e(d1, 1)$ 、 $e(d1, 2)$ 、 $e(d1, 3)$ であったとすれば、以下に示す3組のデータ対によって、単位区間 $d1$ の音響信号を表現することができる。

$$\begin{aligned} &n(d1, 1), e(d1, 1) \\ &n(d1, 2), e(d1, 2) \\ &n(d1, 3), e(d1, 3) \end{aligned}$$

【0050】

以上、単位区間 $d1$ についての処理について説明したが、単位区間 $d2 \sim d5$ についても、それぞれ別個に同様の処理が行われ、代表符号コードおよびその強度を示すデータが得られることになる。たとえば、単位区間 $d2$ については、

$$\begin{aligned} &n(d2, 1), e(d2, 1) \\ &n(d2, 2), e(d2, 2) \\ &n(d2, 3), e(d2, 3) \end{aligned}$$

なる3組のデータ対が得られる。このようにして単位区間ごとに得られたデータによって、原音響信号を符号化することができる。

【0051】

図2は、上述の方法による符号化の概念図である。図2(a)には、図1(a)と同様に、

10

20

30

40

50

原音響信号について5つの単位区間d 1～d 5を設定した状態が示されており、図2(b)には、単位区間ごとに得られた符号データが音符の形式で示されている。この例では、個々の単位区間ごとに3個の代表符号コードを抽出しており(P=3)、これら代表符号コードに関するデータを3つのトラックT 1～T 3に分けて収容するようにしている。たとえば、単位区間d 1について抽出された代表符号コードn(d 1, 1), n(d 1, 2), n(d 1, 3)は、それぞれトラックT 1, T 2, T 3に収容されている。もっとも、図2(b)は、上述の方法によって得られる符号データを音符の形式で示した概念図であり、実際には、各音符にはそれぞれ強度に関するデータが付加されている。たとえば、トラックT 1には、ノートナンバーn(d 1, 1), n(d 2, 1), n(d 3, 1)…なる音階を示すデータとともに、e(d 1, 1), e(d 2, 1), e(d 3, 1)…なる強度を示すデータが収容されることになる。

10

【0052】

なお、ここで採用する符号化の形式としては、必ずしもMIDI形式を採用する必要はないが、この種の符号化形式としてはMIDI形式が最も普及しているため、実用上はMIDI形式の符号データを用いるのが最も好ましい。MIDI形式では、「ノートオン」データもしくは「ノートオフ」データが、「デルタタイム」データを介在させながら存在する。「ノートオン」データは、特定のノートナンバーNとベロシティーVとを指定して特定の音の演奏開始を指示するデータであり、「ノートオフ」データは、特定のノートナンバーNとベロシティーVとを指定して特定の音の演奏終了を指示するデータである。また、「デルタタイム」データは、所定の時間間隔を示すデータである。ベロシティーVは、たとえば、ピアノの鍵盤などを押し下げる速度(ノートオン時のベロシティー)および鍵盤から指を離す速度(ノートオフ時のベロシティー)を示すパラメータであり、特定の音の演奏開始操作もしくは演奏終了操作の強さを示すことになる。

20

【0053】

前述の方法では、第i番目の単位区間d iについて、代表符号コードとしてP個のノートナンバーn(d i, 1), n(d i, 2), …, n(d i, P)が得られ、このそれぞれについて実効強度e(d i, 1), e(d i, 2), …, e(d i, P)が得られる。そこで、次のような手法により、MIDI形式の符号データを作成することができる。まず、「ノートオン」データもしくは「ノートオフ」データの中で記述するノートナンバーNとしては、得られたノートナンバーn(d i, 1), n(d i, 2), …, n(d i, P)をそのまま用いていけばよい。一方、「ノートオン」データもしくは「ノートオフ」データの中で記述するベロシティーVとしては、得られた実効強度e(d i, 1), e(d i, 2), …, e(d i, P)を、値が0～1の範囲となるように規格化し、この規格化後の実効強度Eの平方根に、たとえば127を乗じた値を用いるようにする。すなわち、実効強度Eについての最大値をE_{max}とした場合、

30

$$V = \sqrt{(E / E_{\max})} \cdot 127$$

なる演算で求まる値Vをベロシティーとして用いる。あるいは対数をとって、

$$V = \log(E / E_{\max}) \cdot 127 + 127$$

(ただし、V<0の場合はV=0とする)

なる演算で求まる値Vをベロシティーとして用いてもよい。また、「デルタタイム」データは、各単位区間の長さに応じて設定すればよい。

40

【0054】

結局、上述した実施形態では、3トラックからなるMIDI符号データが得られることになる。このMIDI符号データを3台のMIDI音源を用いて再生すれば、6チャンネルのステレオ再生音として音響信号が再生される。

【0055】

上述した図1および図2を用いた説明では、非常に単純な区間設定例を示したが、実際には、このような区間設定に基いて符号化を行った場合、再生時に、境界となる時刻において音の不連続が発生しやすい。したがって、実用上は、隣接する単位区間が時間軸上で部分的に重複するような区間設定を行うのが好ましい。

50

【0056】

図3(a)は、このように部分的に重複する区間設定を行った例である。図示されている単位区間d1～d4は、いずれも部分的に重なっており、このような区間設定に基いて前述の処理を行うと、図3(b)の概念図に示されているような符号化が行われることになる。この例では、それぞれの単位区間の中心を基準位置として、各音符をそれぞれの基準位置に配置しているが、単位区間に対する相対的な基準位置は、必ずしも中心に設定する必要はない。図3(b)に示す概念図を図2(b)に示す概念図と比較すると、音符の密度が高まっていることがわかる。このように重複した区間設定を行うと、作成される符号データの数は増加することになるが、再生時に音の不連続が生じない自然な符号化が可能になる。

10

【0057】

<<< §2. 人間の声に対する音高シフト >>>

さて、§1で述べた技術を利用すれば、任意の音響信号をMIDIデータなどのデジタル符号データに変換することができるので、楽器の演奏音に限らず、人間の話し声や歌声を符号化することが可能であり、人間の声を、五線譜上の音符として表現することも可能である。もちろん、MIDI規格は、もともと楽器演奏の操作を記述するための符号化規格であるため、個々の符号は、基本的に、特定周波数の音（特定の鍵盤の音）が特定時間だけ持続する（特定時間だけ鳴る）ことを示しているにすぎない。したがって、符号化したMIDIデータを再生、すなわち、所定の音源を用いて演奏しても、元の人間の声がそのまま再生されるわけではない。ただ、楽器を使って人間のしゃべる声に似せた演奏を行うことができるので、エンターテインメントとして様々な利用形態が広がることになる。

20

【0058】

そのような利用形態を考えた場合、MIDIデータとして表現された人間の声について、音高を変化させたいという要求が生まれるのは当然である。たとえば、MIDIデータとして表現された声についてその抑揚を変化させたいという場合や、MIDIデータとして表現された声に節をつけて歌声を作りたいという場合、当該MIDIデータに対して音高をシフトする処理が必要になる。このような音高シフトは、音楽でいう移調演奏に相当するものであり、基本的には、一律に周波数を上げ下げする処理によって行うことができる。

【0059】

たとえば、図4(a)に示すような和音を考えてみよう。この和音は、3つの音符n1, n2, n3から構成されている。本願明細書では、便宜上、個々の音符n1, n2, n3についてのノートナンバーも、同じ記号n1, n2, n3で表すことにする。ここで、ノートナンバーnは、MIDI規格上で定義された0～127までの数値であり、半音の差をもった鍵盤（ピアノの白鍵および黒鍵）の番号を示す数字である。任意のノートナンバーnをもつ音符に対して、ノートナンバー(n+1)をもつ音符は半音だけ高い音符を示し、ノートナンバー(n-1)をもつ音符は半音だけ低い音符を示している。また、本願明細書では、ノートナンバーn1, n2, n3に対応する周波数をそれぞれf1, f2, f3と表すことにする。図4(a)に括弧書きで示すf1, f2, f3は、それぞれ各音符n1, n2, n3に対応する周波数を示している。

30

40

【0060】

ノートナンバーnは、周波数fの対数値（2を底とする対数）に比例しており、周波数が2倍になるたびに、ノートナンバーは12だけ増加し、周波数が1/2倍になるたびに、ノートナンバーは12だけ減少する。西洋音楽でいう「1オクターブの音程」は、周波数が2倍となる関係にある音の差を意味しており、ノートナンバーでは12の差に対応する。いわゆる「ドレミファソラシド」の音階において、「ミ」と「ファ」の間、および「シ」と「ド」の間のみが半音、その他は全音となっているため、12-半音の差が1オクターブに相当する。

【0061】

図4(b)は、図4(a)に示す3つの音符n1, n2, n3を、それぞれ1オクターブず

50

つ上げることにより得られる新たな音符 n_1^* , n_2^* , n_3^* を示している。すなわち、図 4 (a) に示す 3 つの音符 n_1 , n_2 , n_3 (いずれも「ラ」の音に対応する音符) が、図 4 (b) では、1 オクターブ高い (12 - 半音だけ高い) 3 つの音符 n_1^* , n_2^* , n_3^* (これらも「ラ」の音に対応する音符) に置き換えられている。図に括弧書きで示すとおり、各音符 n_1^* , n_2^* , n_3^* に対応する周波数は f_1^* , f_2^* , f_3^* である。ここで、ノートナンバーの関係は、 $n_1^* = n_1 + 12$ 、 $n_2^* = n_2 + 12$ 、 $n_3^* = n_3 + 12$ であるが、周波数の関係は、 $f_1^* = f_1 \times 2$ 、 $f_2^* = f_2 \times 2$ 、 $f_3^* = f_3 \times 2$ となる。

【0062】

結局、図 4 (b) に示す和音は、図 4 (a) に示す和音を構成する音符を一律「1 オクターブ」上げる (すなわち、周波数 f としては 2 倍にする、ノートナンバー n としては 12 を加算する) ことによって得られる。

【0063】

このように、一般的には、任意の和音について、音高を上げ下げする処理とは、当該和音を構成する個々の音符のノートナンバーに対して、一律に所定のオフセット値を加減算することによって行われる。上例のように、1 オクターブだけ上げる音高シフトを行うのであれば、すべての音符のノートナンバー n に対して、1 オクターブの音程に相当するオフセット値「12」を加える加算処理を行えばよい。

【0064】

したがって、何らかの楽曲を表す M I D I データに含まれているすべての符号について、そのノートナンバー値を 12 だけ増加させる変換処理を行えば、変換後の M I D I データは、元の M I D I データに対して 1 オクターブ分の音高シフト処理を施した移調演奏用のデータということになる。そして、一般的には、そのような移調演奏用の M I D I データを再生 (所定の音源を用いて演奏) したとしても、不自然と思われるほどの違和感が生じない。これは、M I D I 規格やピアノ鍵盤では平均律音階を採用しているためである。西洋音楽の音階「ドレミファソラシド」は、紀元前のピタゴラス音階が原点で、当初はこれを改良した純正律音階 (構成音の周波数比が全て整数になる) が使用されていた。しかし、移調演奏を行うと、協和音が不協和音になるなど、メロディと和声の関係が崩れてしまうという問題があり、17 世紀に J S バッハらにより平均律音階が発明された。これは 1 オクターブ内の構成音を 12 音に等比級数で分割する方法で、構成音の周波数比が全て $2^{-1/12}$ を基本とした端数をもつ実数比になるため、協和音でも若干の唸り音が生じるという欠点はあるが、いかなる移調を施しても (半音単位にノートナンバーを上下させても)、メロディと和声の関係が常に維持されるという利点があり、以降の音楽は殆ど平均律音階に基づいて作曲されている。

【0065】

ところが、既に述べたとおり、本願発明者が行った実験によると、人間の声を表現した M I D I データについて、同様の方法で音高シフトを行うと、音声の明瞭性は維持されるものの、声質が破壊され、違和感のある奇怪な声になってしまうことが判明した。具体的には、プロのアナウンサーによって伝えられた天気予報のニュース番組を録音し、§ 1 で述べた方法で符号化して M I D I データを作成し、当該 M I D I データに含まれているすべての符号について、そのノートナンバー値を 12 だけ増加させる変換処理を行い、変換後の M I D I データを再生したところ、個々の単語の判別は可能であるものの、いわゆるボイスチェンジャーを通したような違和感のある不自然な声になってしまった。

【0066】

本願発明者は、そのような結果が得られる理由を探求するために種々の実験を行った結果、人の声に含まれているフォルマント成分に対して行った音高シフト処理が、違和感や不自然さの原因になっていることを見出した。以下、この原因について、具体例を挙げて説明する。

【0067】

一般に、楽器音や人間が発声する音声には、基本周波数成分と、その整数倍の周波数を

もつ倍音成分が離散的に含まれており、特に音声の場合は、倍音成分の中で強度が前後の倍音に比べ大きなピークをもつ成分がみられ、これらの特別な倍音はフォルマント成分と命名されている。（これに対し、楽器音の場合は周波数が高い倍音ほど強度が単調に小さくなる傾向を示すのが一般的である。）フォルマント成分は基本周波数を F_0 と命名し、周波数が低い方から順に F_1 、 F_2 、 F_3 ・・・と命名される。図5は、人間の声に含まれるフォルマント成分を例示する強度グラフである。このグラフは、横軸に周波数 f をとり、縦軸に実効強度 E をとった、人間の音声スペクトルのピーク位置を示すグラフである。人間の声の周波数スペクトルは、かなり広い周波数域にわたって分布するが、ところどころの周波数位置にピークが存在する。図5は、このようなピークのみを抜き出して示したグラフであり、個々の周波数位置に示された垂直のバーは、当該周波数位置に当該バーの長さに相当する実効強度をもつピークが存在することを示している。ここで、太線のバーは、特に強度の大きなピークを示している。すなわち、図示の例の場合、周波数 F_0 の位置に最大ピークが見られ、以下、周波数 F_1 、 F_2 、 F_3 の位置のピークが続く。

【0068】

ここで、基本周波数と一致する最小周波数に位置するフォルマントの周波数は F_0 と呼ばれており、この基本周波数 F_0 をもった成分は、人間の声の最も主要な成分ということになる。これに対して、その他の太線のバーで示されている周波数 F_1 、 F_2 、 F_3 をもつ成分は、低い方から順に、それぞれ第1フォルマント周波数 F_1 、第2フォルマント周波数 F_2 、第3フォルマント周波数 F_3 と呼ばれる。これら各フォルマント周波数 F_1 、 F_2 、 F_3 は、基本周波数 F_0 の整数倍の位置に現れることが知られている。基本周波数 F_0 は、人間の声帯の本来の振動周波数に対応するものであるが、この基本周波数 F_0 の整数倍をとる倍音系列のうち、声道共鳴によっていくつかの倍音が顕著なエネルギーピークをもつ。そして、このようにエネルギーピークをもつに至った倍音成分がフォルマントを形成するとされている。

【0069】

図5に示す例の場合、たとえば、基本周波数 $F_0 = 100\text{ Hz}$ （声帯の振動周波数）とすると、他の細線のバーや太線のバーの位置は、いずれもその整数倍である 200 Hz 、 300 Hz 、...となっており、これらの倍音成分のうち、声道共鳴によって顕著なピークをもつに至った周波数 $F_1 (= 600\text{ Hz})$ 、 $F_2 (= 1200\text{ Hz})$ 、 $F_3 (= 1600\text{ Hz})$ の成分が、フォルマントを構成することになる。

【0070】

もちろん、このようなフォルマントの特性には個人差があり、どの倍音成分がフォルマントになるか、第3フォルマント周波数 F_3 以上の高次フォルマントが生じるか否か、といった条件は、発声主である人によって千差万別である。もちろん、しゃべる言語によっても、その特性に変化が生じる。しかし、 F_0 以外の高次のフォルマント周波数は言語ごとに母音ごとにほぼ一定の範囲に収まっているため、音声によるコミュニケーションが成立する。言語習得の過程において自ら発声する母国語の各母音のフォルマント周波数を所定の範囲に収めるように声道の筋肉が親や教師より訓練されている。（鳥も親鳥より同様な訓練がなされていることも知られている。）現在のところ、母音を発音したときのフォルマント特性についてはある程度の解析がなされているものの、子音を発音したときのフォルマント特性については未解明な部分が多い。ただ、いずれにせよ、このフォルマント成分が、人間の声の質に大きな影響を与える要素であることは事実である。

【0071】

本願発明者は、このようなフォルマントによって特徴づけられる人間の声において、音高を変えるような発声を試みた場合、基本周波数 F_0 は上下するものの、フォルマント周波数 F_1 、 F_2 、 F_3 、...には、大きな変化が生じない傾向があることを実験により確認した。具体的には、同一人物に、同一の母音を、標準的な音高、高めの音高、低めの音高の3通りで発声してもらい、そのスペクトルをとり、図5のグラフに例示するようなピーク位置を求める実験を行った。図6(a)、(b)、(c)に示すグラフは、このような実験の一例を示す強度グラフである。すなわち、図6(a)は、標準的な音高で発声を行った音

声スペクトルの強度グラフ、図 6 (b) は、音高を高めて発声を行った場合の強度グラフ、図 6 (c) は、音高を低めて発声を行った場合の強度グラフである。

【0072】

図 6 (a) , (b) , (c) において、破線は、基本周波数 F_0 の倍音位置を示している。たとえば、図 6 (a) における基本周波数 $F_0 = 100 \text{ Hz}$ だとすると、図 6 (a) に示す破線は 100 Hz の間隔で引かれていることになり、この例の場合、第 1 フォルマント周波数 $F_1 = 600 \text{ Hz}$ 、第 2 フォルマント周波数 $F_2 = 1200 \text{ Hz}$ 、第 3 フォルマント周波数 $F_3 = 1600 \text{ Hz}$ である。一方、図 6 (b) における基本周波数 $F_0 = 200 \text{ Hz}$ だとすると、図 6 (b) に示す破線は 200 Hz の間隔で引かれていることになり、図 6 (c) における基本周波数 $F_0 = 50 \text{ Hz}$ だとすると、図 6 (c) に示す破線は 50 Hz の間隔で引

10

【0073】

このように、いずれの場合も、ピーク強度が得られる周波数（細線のバーや太線のバーで示す周波数）が、基本周波数 F_0 の倍音位置（破線の位置）になる点に変わりはない。ただ、これらのピーク強度が得られる周波数（倍音系列の周波数）のうち、顕著な強度値をもつフォルマント周波数 F_1 、 F_2 、 F_3 （太線のバー）に着目すると、基本周波数 F_0 の変化にもかかわらず、全く変化していないことがわかる。

【0074】

たとえば、図 6 (b) の場合、基本周波数 F_0 は、 100 Hz から 200 Hz にシフトしている（発声者が高い声を出そうとしたため、声帯の振動周波数が高くなっている）。ところが、フォルマント周波数 F_1 、 F_2 、 F_3 の位置は、 600 Hz 、 1200 Hz 、 1600 Hz であり、図 6 (a) の場合と全く変わりはない。同様に、図 6 (c) の場合、基本周波数 F_0 は、 100 Hz から 50 Hz にシフトしている（発声者が低い声を出そうとしたため、声帯の振動周波数が低くなっている）。ところが、フォルマント周波数 F_1 、 F_2 、 F_3 の位置は、 600 Hz 、 1200 Hz 、 1600 Hz であり、図 6 (a) の場合と全く変わりはない。

20

【0075】

この図 6 に示す例は、説明の便宜上、最も単純な典型例を示すものであり、実際の実験では、必ずしもこのような結果が得られるわけではない。上述したとおり、フォルマント特性には個人差があり、基本周波数 F_0 が変わることによって、新たなフォルマント成分が追加されたり、逆に一部のフォルマント成分が消失したりするケースもある。

30

【0076】

また、図 6 に示す例では、図 6 (a) に示す基本周波数 F_0 に対して、図 6 (b) に示す基本周波数 F_0 はたまたま「正確に 2 倍」となっており、図 6 (c) に示す基本周波数 F_0 はたまたま「正確に $1/2$ 倍」となっているため、破線で示す倍音位置に関して、図 6 (a) , (b) , (c) 間で整合性がとれ、フォルマント周波数 F_1 、 F_2 、 F_3 の位置が、3 つの図で完全に一致している。しかしながら、実際の実験では、基本周波数 F_0 の値は、発声者の「標準的な声」、「高めの声」、「低めの声」といった感覚で決められる値であるから、当然ながら、「正確に 2 倍」や「正確に $1/2$ 倍」といった値にはならない。したがって、図に破線で示す倍音位置も、グラフ間で完全に整合性がとれるわけではない。

40

【0077】

たとえば、図 6 (b) における基本周波数 F_0 が、 190 Hz であったとすると、図 6 (b) の破線の間隔（倍音位置）は 190 Hz おきになるので、フォルマント周波数 F_1 、 F_2 、 F_3 の位置は、理論的には、 570 Hz 、 1140 Hz 、 1520 Hz となり、図 6 (a) におけるフォルマント周波数 F_1 、 F_2 、 F_3 の位置からは若干ずれることになる。

【0078】

このように、実際に行った実験では、図 6 に示すような典型的な結果が得られるわけではない。しかしながら、同じ発声者が同じ母音を発音しながら、音高を上昇あるいは下降させると、基本周波数 F_0 は、それに連動して上下し、新たな基本周波数 F_0 の倍音系列を含む音が再構成されることは確かである。しかも、基本周波数 F_0 をシフトさせても、

50

各フォルマント周波数 F_1 , F_2 , F_3 の位置はそれほど大きくシフトせず、「高めの声」を出した場合でも、あるいは「低めの声」を出した場合でも、「標準的な声」を出した場合の各フォルマント周波数 F_1 , F_2 , F_3 の近傍位置にそのままとどまることが確認できた。これは、基本周波数 F_0 が声帯の振動によって定まる周波数であるのに対して、各フォルマント周波数 F_1 , F_2 , F_3 は、声道での共鳴に依存して定まる周波数であるためと考えられる。

【0079】

このような事実を踏まえれば、図 6 (a) に示すような基本周波数 F_0 およびフォルマント周波数 F_1 , F_2 , F_3 をもった人間の声を符号化した M I D I データに対して、音高を 1 オクターブだけ上げる音高シフト処理を行うのであれば、基本周波数 F_0 の音については 2 倍の周波数に変更するが、フォルマント周波数 F_1 , F_2 , F_3 の音については周波数変更を一切行わない、という処理を行うのが好ましいことがわかる。同様に、音高を 1 オクターブだけ下げる音高シフト処理を行うのであれば、基本周波数 F_0 の音については $1/2$ 倍の周波数に変更するが、フォルマント周波数 F_1 , F_2 , F_3 の音については周波数変更を一切行わない、という処理を行うのが好ましいことがわかる。

【0080】

結局、人間の声を表現した M I D I データについて、一般の楽曲を表現した M I D I データについての音高シフト方法をそのまま適用すると、ボイスチェンジャーを通したような違和感のある不自然な声になる理由は、すべての周波数の音に対して、一律に音高シフト処理を施すと、本来は音高シフトすべきではないフォルマント成分に対しても音高シフト処理が行われてしまうためであることがわかる。したがって、人間の声に対する音高シフトを行う際には、「フォルマント成分の絶対周波数値はほぼ一定」という人間の声音に固有の特徴を維持したまま、音高を変更する処理を行う必要がある。

【0081】

<<< § 3. 本発明の基本概念 >>>

§ 2 で述べた人間の声の特性を考慮すれば、人間の声を表現した M I D I データについての音高シフトを行う際には、基本周波数 F_0 をもつ音には音高のシフト処理（周波数／ノートナンバーの増減処理）を行い、フォルマント周波数 F_1 , F_2 , F_3 をもつ音には音高のシフト処理を行わない、という分別処理を行うとよいことがわかる。しかしながら、実際にそのような分別処理を行うことは非常に困難である。たとえば、図 6 (a) に示す例の場合、基本周波数 F_0 は、最も大きな強度ピークをとる周波数であり、かつ、複数の強度ピークに対応する周波数の中で最も小さな周波数となっている。このような典型的な例の場合、基本周波数 F_0 とフォルマント周波数 F_1 , F_2 , F_3 とを弁別することは容易であり、基本周波数 F_0 をもつ音に対してのみシフト処理を行うことができる。

【0082】

しかしながら、図 6 (a) に示す例は、説明の便宜のために示した理想的な典型例であり、実際の人間の声に対して、§ 1 で述べた符号化方法を適用することによって得られる M I D I データの場合、必ずしもこの典型例のような結果が得られるわけではない。したがって、実用上は、基本周波数 F_0 とフォルマント周波数 F_1 , F_2 , F_3 とを簡単に弁別できるものではなく、実際に分別を行うとなると、高度な音声信号解析技術が必要になる。上述したとおり、フォルマント特性には個人差があり、しかも現在の技術では、子音についてのフォルマントに関する説明は十分になされていないため、一般的な人間の話し言葉や歌声について、最も大きなピーク、あるいは、最も周波数の小さなピーク、といった単純な方法で基本周波数 F_0 を認識することは困難である。

【0083】

そこで、本願発明者は、人間の声を表現した符号化音声データに対して、できるだけ簡便な方法により所望の音高シフト処理を施し、違和感のない自然な音声再生が可能な符号化音声データを得る方法を探るうちに、次のような着想を得ることができた。

【0084】

いま、図 7 (a) に示すような重み関数 $W(f)$ を定義してみる。この図 7 (a) のグラフ

は、横軸に周波数 f （線形スケール）、縦軸に関数値 $W(f)$ をとり、周波数 f 軸上の所定区間 $f_a \sim f_b$ において関数値 $W(f)$ が周波数 f の増加に従って単調減少する関数を示している。より具体的に言えば、この重み関数 $W(f)$ は、周波数 f 軸上の第1設定値 f_a および第2設定値 f_b ($f_b > f_a$) について、 $f \leq f_a$ の区間は、 $W(f) = 1$ 、 $f_a < f < f_b$ の区間は、 $1 > W(f) > 0$ （但し、 $W(f)$ は f の増加に従って単調減少）、 $f \geq f_b$ の区間は、 $W(f) = 0$ となる関数である。

【0085】

さて、ここで、MIDI データに対して、オフセット値 α に応じた音高シフト処理を行う際に、この MIDI データに含まれている個々の符号（音符）について、当該符号が示す周波数 f を $\alpha \cdot W(f)$ に応じた値だけ増減することにより新たな周波数 f' を求め、当該符号を、それが示す周波数 f を f' に変更した新たな符号に置き換える処理を行うことにする。

10

【0086】

たとえば、音高を高める場合は正、低める場合は負のオフセット値 α を設定し、もとの周波数 f に対して、所定の係数 k ($k > 1$) を用いた式 $f' = f \cdot k^{\alpha \cdot W(f)}$ により新たな周波数 f' を求めればよい。

【0087】

具体的には、 $\alpha = 1$ に設定した場合、 $f' = f \cdot k^{W(f)}$ なる演算により新たな周波数 f' が求まる。この場合、図7のグラフにおいて、第1設定値 f_a 以下の周波数 f に対しては、 $W(f) = 1$ なる重み関数値が与えられるので、 $f' = k \cdot f$ となり、周波数は k 倍にシフトする（たとえば、 $k = 2$ に設定しておけば、音程は1オクターブ上がることになる）。これに対して、第2設定値 f_b 以上の周波数 f に対しては、 $W(f) = 0$ なる重み関数値が与えられるので、 $f' = f$ となり、周波数は変化しない。第1設定値 f_a を超え、第2設定値 f_b 未満の周波数 f に対しては、その中間的な量（ k 倍～1倍）の周波数シフトが行われることになる。

20

【0088】

また、 $\alpha = -1$ に設定した場合も、図7のグラフにおいて、第1設定値 f_a 以下の周波数 f に対しては、 $W(f) = 1$ なる重み関数値が与えられるので、 $f' = f / k$ となり、周波数は $1/k$ 倍にシフトする（たとえば、 $k = 2$ に設定しておけば、音程は1オクターブ下がることになる）。これに対して、第2設定値 f_b 以上の周波数 f に対しては、 $W(f) = 0$ なる重み関数値が与えられるので、 $f' = f$ となり、周波数は変化しない。第1設定値 f_a を超え、第2設定値 f_b 未満の周波数 f に対しては、その中間的な量（ $1/k$ 倍～1倍）の周波数シフトが行われることになる。

30

【0089】

このように、周波数 f の増加に従って単調減少する重み関数 $W(f)$ を用いて、符号が示す周波数 f を $\alpha \cdot W(f)$ に応じた値だけ増減するようにすると、より高音の符号ほど音高シフト量が減少することになる。このため、人間の発生する音声についての音高シフトを行う場合、基本周波数成分の音高シフト量に比べて、当該基本周波数成分の整数倍の周波数をもつフォルマント成分についての音高シフト量が小さくなり、「フォルマント成分の絶対周波数値はほぼ一定」という人間の声の特徴をできるだけ維持したまま、音高を変更することができる。

40

【0090】

ここで、第1設定値 f_a は、平均的な音声についての基本周波数成分（F0：0次フォルマント成分）が含まれると予想される周波数領域の上限に設定し、第2設定値 f_b は、聴取可能な最高次フォルマント $F^{\max}$ 成分が含まれると予想される周波数領域の上限に設定するのが好ましい。

【0091】

具体的には、第1設定値 f_a は、音高シフト処理の対象となる音声が男性の声の場合、100Hz～200Hzの範囲内に設定し、音高シフト処理の対象となる音声が女性の声の場合、200Hz～400Hzの範囲内に設定するのが好ましい。これは、一般的な男

50

性の声の場合、その基本周波数成分が含まれると予想される周波数領域の上限は、概ね 100 Hz ～ 200 Hz あたりとなり、一般的な女性の声の場合、その基本周波数成分が含まれると予想される周波数領域の上限は、概ね 200 Hz ～ 400 Hz あたりとなるためである。

【0092】

一方、第2設定値 f_b は、音高シフト処理の対象となる音声が男性の声の場合、3 kHz ～ 6 kHz の範囲内に設定し、音高シフト処理の対象となる音声が女性の声の場合、4 kHz ～ 8 kHz の範囲内に設定するのが好ましい。これは、一般的な男性の声の場合、人間が聴取可能な最高次フォルマント $F^{\max}$ 成分が含まれると予想される周波数領域の上限は、概ね 3 kHz ～ 6 kHz あたりとなり、一般的な女性の声の場合、人間が聴取可能な最高次フォルマント $F^{\max}$ 成分が含まれると予想される周波数領域の上限は、概ね 4 kHz ～ 8 kHz あたりとなるためである。

10

【0093】

既に述べたとおり、基本周波数 F_0 や各フォルマント周波数 F_1 、 F_2 、 F_3 が、周波数軸上のどこに位置するかという事項は、個人差や言語差によって大きく左右される事項であり、一概に決定することはできない。しかしながら、図7(a)に示すような重み関数 $W(f)$ を設定し、周波数 f を $\alpha \cdot W(f)$ に応じた値だけ増減するようにすれば、「より高音の符号ほど音高シフト量が減少する」という一般的な傾向をもった音高シフト処理を行うことができるので、たとえ個人差や言語差によるバリエーションがあったとしても、かなり広いバリエーションをカバーした概括的な適用が可能になり、概ね良好な処理結果を得ることができる。

20

【0094】

たとえば、ある人物については、その声の基本周波数 F_0 が第1設定値 f_a を超えている場合もある。そのような場合でも、基本周波数 F_0 が第2設定値 f_a に達しない限り ($f_a = 3 \text{ kHz}$ 程度に設定しておけば、通常の話し言葉では、基本周波数 F_0 が第2設定値 f_a に達することは考えにくい)、関数値 $W(f)$ は0にはならないので、オフセット値 α に応じたシフト量で、基本周波数 F_0 に対する音高シフト処理が実行される。

【0095】

また、多くの場合、いくつかのフォルマント周波数は、第1設定値 f_a ～第2設定値 f_b の範囲内に位置することになるであろう。それでも、これらフォルマント周波数に対する音高シフト量は、基本周波数 F_0 に対する音高シフト量よりも必ず小さくなるので、音高シフト処理後の音声をできるだけ違和感のない自然な音声に近づけるという効果は少なからず得られることになる。

30

【0096】

実際、本願発明者が試した多くの事例では、上述したように、第1設定値 f_a を、男性の声の場合には、100 Hz ～ 200 Hz の範囲内に設定し、女性の声の場合には、200 Hz ～ 400 Hz の範囲内に設定し、第2設定値 f_b を、男性の声の場合には、3 kHz ～ 6 kHz の範囲内に設定し、女性の声の場合には、4 kHz ～ 8 kHz の範囲内に設定すれば、ほぼすべてのケースについて、違和感のない自然な音声再生が可能な符号化音声データを得るという本発明に特有の作用効果を得ることができた。

40

【0097】

<<< § 4. ノートナンバーの増減による音高シフト >>>

§ 3では、本発明の基本概念を、「周波数をシフトする」という物理的な事象の観点から述べたが、MIDI規格などの符号化音声データの場合、音高は周波数 f ではなく、ノートナンバー n によって表現される。したがって、実用上は、「周波数をシフトする」処理の代わりに、「ノートナンバーをシフトする」処理を行う必要がある。そこで、以下、このような実用上の観点から本発明の基本概念を説明する。

【0098】

既に述べたとおり、ノートナンバー n は、周波数 f の対数値に比例しており、周波数が2倍になるたびに、ノートナンバーは12だけ増加し、周波数が1/2倍になるたびに、

50

ノートナンバーは12だけ減少する関係にある。したがって、周波数 f の増減によって音高シフトを行う代わりに、ノートナンバーの増減によって音高シフトを行う場合、重み関数も、周波数 f の関数 $W(f)$ の代わりに、ノートナンバー n の関数 $W(n)$ を用いる必要がある。

【0099】

図7(b)は、このようなノートナンバー n についての重み関数 $W(n)$ の一例を示すグラフである。この図7(b)のグラフは、横軸がノートナンバー n 、すなわち、周波数 f についての対数スケールとなっているため、図7(a)のグラフと形状は異なっているが、実際には、関数 $W(f)$ と関数 $W(n)$ とは等価な関数である。関数 $W(n)$ は、ノートナンバー n の座標軸上に設定された第1設定値 n_a 、第2設定値 n_b について、 $n \leq n_a$ の区間は、 $W(n) = 1$ 、 $n_a < n < n_b$ の区間は、 $1 > W(n) > 0$ (但し、 $W(n)$ は n の増加に従って単調減少)、 $n \geq n_b$ の区間は、 $W(n) = 0$ となる関数である。特に、この図7(b)に示す例の場合、 $n_a < n < n_b$ の区間において、 $W(n)$ の値は、ノートナンバー n に反比例する値となっている。

10

【0100】

ここで、図7(b)に示すノートナンバーについての第1設定値 n_a 、第2設定値 n_b は、それぞれ図7(a)に示す周波数についての第1設定値 f_a 、第2設定値 f_b に対応するものである。別言すれば、周波数 f_a はノートナンバー n_a に対応し、周波数 f_b はノートナンバー n_b に対応する。結局、図7(b)に示す重み関数 $W(n)$ のグラフは、図7(a)に示す重み関数 $W(f)$ のグラフの周波数軸を対数スケールに修正したものに相当する。

20

【0101】

なお、図7(b)に示す例において、 $n_a < n < n_b$ の区間についての $W(n)$ の値が、ノートナンバー n に反比例する値となるようにしているのは、演算負担を軽減するための配慮である。この図7(b)に示すような重み関数 $W(n)$ を用いるようにすれば、 $n \leq n_a$ の区間は、 $W(n) = 1$ 、 $n \geq n_b$ の区間は、 $W(n) = 0$ 、そして $n_a < n < n_b$ の区間は、 $W(n) = k/n$ (k は比例定数) となり、関数値 $W(n)$ を求める演算負担は極めて軽くなる。もちろん、 $W(n)$ の値は、 $n_a < n < n_b$ の区間について、 n の増加に従って単調減少すればよいので、必ずしも反比例するような値に設定する必要はない。

【0102】

MIDIデータに対する音高シフト処理は、§2で述べたとおり、MIDI符号に含まれるノートナンバー n に対して、オフセット値 α を加算もしくは減算することによって行われる。すなわち、1オクターブだけ上げる音高シフトを行うのであれば、ノートナンバー n に対して、1オクターブの音程に相当するオフセット値「12」を加える加算処理を行い、1オクターブだけ下げる音高シフトを行うのであれば、ノートナンバー n に対して、1オクターブの音程に相当するオフセット値「12」を減じる減算処理を行えばよい。

30

【0103】

そこで、ここでは、具体例を参照しながら、従来の方法による音高シフト処理と、本発明による音高シフト処理との相違を説明しよう。いま、図8(a)に示すように、変換前のMIDIデータに、時間軸上で同一期間を占める3つの符号が含まれていたものとしよう。図8(a)では、この3つのMIDI符号を音符により表している。実際のMIDI符号には、前述したとおり、「ノートナンバー」、「ベロシティ」、「ノートオン」、「ノートオフ」、「デルタタイム」などのデータが含まれるが、音高シフト処理で変更されるデータは、周波数を示す「ノートナンバー」のみであるから、以下、便宜上、MIDI符号を音符で表現した説明を行うことにする。

40

【0104】

図8(a)に示す例の場合、音高シフト処理の対象となる音は、3つの音符 n_1 、 n_2 、 n_3 の和音として構成される音である。これらの各音符 n_1 、 n_2 、 n_3 は、それぞれ同じ記号で示すノートナンバー n_1 、 n_2 、 n_3 をもった音を示し、これらのノートナンバーは、それぞれ括弧書きで示す周波数 f_1 、 f_2 、 f_3 に対応するものとする。

50

【0105】

§1で述べた方法で、人間の声をMIDIデータによって表すと、図3(b)に示す例のように、複数のトラックにそれぞれ符号が得られることになり、これら複数トラックに分布している符号を五線譜上に音符として表現すると、図8(a)に例示するような和音が形成されることになる。要するに、人間の声を表したMIDIデータは、音の持続時間が時間軸上で同一期間を占め、互いに異なる周波数を示す複数の符号によって構成され、五線譜上では和音の形式で表されることになる。

【0106】

さて、このような和音に対して、従来の方法による音高シフト処理を行う場合、§2で述べたとおり、和音を構成する個々の音符に対して、一律、同じシフト量だけノートナンバーを増減する処理が施される。たとえば、1オクターブだけ上げる処理を行う場合、すべての音符 n_1 、 n_2 、 n_3 のノートナンバーに、それぞれ一律して、1オクターブの音程に相当するオフセット値「12」を加える加算処理が行われる。図8(b)は、このような音高シフト処理を施した後に得られる新たな音符 n_1^* 、 n_2^* 、 n_3^* を示している。ここで、各音符 n_1^* 、 n_2^* 、 n_3^* のノートナンバーを同じ記号 n_1^* 、 n_2^* 、 n_3^* で示せば、 $n_1^* = n_1 + 12$ 、 $n_2^* = n_2 + 12$ 、 $n_3^* = n_3 + 12$ が成り立つ。また、シフト処理後の各ノートナンバーに対応する周波数を f_1^* 、 f_2^* 、 f_3^* とすれば、 $f_1^* = f_1 \times 2$ 、 $f_2^* = f_2 \times 2$ 、 $f_3^* = f_3 \times 2$ が成り立ち、いずれの音符についても一律に周波数を2倍にする処理が行われていることになる。

【0107】

このように、すべての音符のノートナンバーについて一律に同じシフト量を増減する音高シフトは、楽器音を想定した一般的な楽曲に対して行う場合は問題ないが、人間の声に対して行くと、基本周波数成分のみならず、フォルマント成分に対しても同じシフト量の音高シフトが行われてしまうため、問題が生じることは既に述べたとおりである。

【0108】

そこで、本発明の方法による音高シフト処理では、オフセット値 α に対して、図7(b)に示す重み関数 $W(n)$ を乗じることにより、ノートナンバーの大きな音符ほど、シフト量が小さくなるような調整を行うことになる。具体的には、「 $n' = n + \alpha \cdot W(n)$ 」なる演算式により新たなノートナンバー n' を求めることになる。

【0109】

図8(c)は、本発明に係る音高シフト処理を施した後に得られる新たな音符 n_1^* 、 n_2^* 、 n_3^* を示している。ここで、各音符 n_1^* 、 n_2^* 、 n_3^* のノートナンバーを同じ記号 n_1^* 、 n_2^* 、 n_3^* で示せば、 $n_1^* = n_1 + 12$ 、 $n_2^* = n_2 + 6$ 、 $n_3^* = n_3 + 3$ となっており、シフト量は音符ごとに異なる。これは、オフセット値 $\alpha = 12$ を与えたとしても、重み関数値が、 $W(n_1) = 1$ 、 $W(n_2) = 0.5$ 、 $W(n_3) = 0.25$ と、ノートナンバーの増加にともなって徐々に減少してゆくため、「 $n' = n + \alpha \cdot W(n)$ 」なる演算式により新たなノートナンバー n' を求めると、シフト量「 $\alpha \cdot W(n)$ 」は、12、6、3と徐々に減少してゆくためである。シフト処理後の各ノートナンバーに対応する周波数を f_1^* 、 f_2^* 、 f_3^* とすれば、 $f_1^* = f_1 \times 2$ 、 $f_2^* = f_2 \times 1.41$ 、 $f_3^* = f_3 \times 1.19$ になる。

【0110】

このように、図8(a)のような和音として表される人間の声のMIDIデータに対して、1オクターブ上げる音高シフト処理を行う場合、本発明による方法を適用すれば、図8(c)のような和音に対応するMIDIデータが得られることになる。この場合、基本周波数成分に対応する可能性の高い音符 n_1 については、1オクターブの音程に相当するオフセット値「12」のシフト量がそのまま加算され、音符 n_1^* に変換されることになるが、フォルマント成分に対応する可能性の高い音符 n_2 、 n_3 についてのシフト量はより小さくなる。このため、人間の声のMIDIデータに対して、できるだけ違和感のない自然な音高シフト処理が可能になる。

【0111】

10

20

30

40

50

図 9 は、1 オクターブ下げる音高シフト処理を示す例である。図 9 (a) に示すような 3 つの音符 n_1 , n_2 , n_3 の和音として構成される音に対して、従来の方法による音高シフト処理を行い、1 オクターブだけ下げる処理を行うと、図 9 (b) に示す音符 n_1^* , n_2^* , n_3^* が得られる。ここで、新たに得られた音符のノートナンバー n_1^* , n_2^* , n_3^* は、 $n_1^* = n_1 - 12$, $n_2^* = n_2 - 12$, $n_3^* = n_3 - 12$ で表される。また、これらの各ノートナンバーに対応する周波数を f_1^* , f_2^* , f_3^* とすれば、 $f_1^* = f_1 / 2$, $f_2^* = f_2 / 2$, $f_3^* = f_3 / 2$ が成り立ち、いずれの音符についても一律に周波数を $1/2$ 倍にする処理が行われていることになる。

【0112】

これに対して、図 9 (c) は、本発明に係る音高シフト処理を施した後に得られる新たな音符 n_1^* , n_2^* , n_3^* を示している。ここで、各音符のノートナンバー n_1^* , n_2^* , n_3^* は、 $n_1^* = n_1 - 12$, $n_2^* = n_2 - 6$, $n_3^* = n_3 - 3$ となっており、やはりシフト量は音符ごとに異なる。シフト処理後の各ノートナンバーに対応する周波数を f_1^* , f_2^* , f_3^* とすれば、 $f_1^* = f_1 / 2$, $f_2^* = f_2 / 1.41$, $f_3^* = f_3 / 1.19$ になる。

【0113】

このように、図 9 (a) のような和音として表される人間の声の M I D I データに対して、1 オクターブ下げる音高シフト処理を行う場合、本発明による方法を適用すれば、図 9 (c) のような和音に対応する M I D I データが得られることになる。この場合、基本周波数成分に対応する可能性の高い音符 n_1 については、1 オクターブの音程に相当するオフセット値「12」のシフト量がそのまま減算され、音符 n_1^* に変換されることになるが、フォルマント成分に対応する可能性の高い音符 n_2 , n_3 についてのシフト量はより小さくなる。このため、人間の声の M I D I データに対して、違和感のない自然な音高シフト処理が可能になる。

【0114】

ところで、和音を構成する複数の音符に対して、それぞれ異なるシフト量をもった音高シフトを行うと、複数の音符がシフト後に同じ音高を占める場合がありうる。M I D I 規格では、このように同じ音高を占める符号が同一チャンネルにおいて、同一時刻に重複して存在することは許されない。ただし、最大 16 種類定義可能なチャンネル番号を変えれば重複が認められる。例えば、ピアノ音をチャンネル 0、バイオリン音をチャンネル 1 に設定すれば、これらの楽器音を同一音高で同一時刻に演奏させることができる。しかし、本願では音声という単一の音色を扱うため、以下は単一のチャンネル番号を使用するという前提で説明する。このように、単一のチャンネル番号を使用する場合、M I D I データに対して、本発明に係る音高シフト処理を行う際には、処理後に同じ音高の符号が重複した場合に、1 つのみを残して他を削除する処理を行うようにすればよい。

【0115】

図 10 は、本発明に係る音高シフト処理を実行すると、複数の音符がシフト後に同じ音高を占めるケースを示す図である。図 10 (a) は、変換前の符号化データに含まれる和音を示し、図 10 (b) は、従来の方法で音高を 1 オクターブ上げる処理を行った結果を示し、図 10 (c) は、本発明の方法で音高を 1 オクターブ上げる処理を行った結果を示す。従来の方法では、図 10 (b) に示すように、全符号のノートナンバーに対して一律に同じオフセット値「12」が加算される。このため、3 つの音符の相互位置関係はそのまま維持され、全体が一群となって平行移動する。ところが、本発明の方法では、3 つの音符はそれぞれシフト量が異なるため、図 10 (c) に示すように、シフト後の 3 つの音符の位置は重なってしまう。

【0116】

すなわち、図 10 (a) に示す音符 n_1 , n_2 , n_3 に対するシフト量は、それぞれ 12, 6, 3 になり、各音符を当該シフト量に応じてシフトすると、本発明に係る音高シフト処理を施した後に得られる新たな音符 n_1^* , n_2^* , n_3^* の位置は、図 10 (c) に示すとおりに重なってしまう。別言すれば、ノートナンバー n_1 , n_2 , n_3 に、それぞれシ

10

20

30

40

50

フト量 12, 6, 3 を加えると、偶然、同じ値になってしまう。

【0117】

MIDI規格では、このように同じ音高をもつ複数の符号が重複することは許されない。そこで、このような場合、図10(c)に重複して示されている新たな音符 $n1^*$, $n2^*$, $n3^*$ のうちの1つのみを残し、その他の2個を削除する重複回避処理を行うようにすれば、重複を許さないMIDI規格のような符号化データに対しても本発明を問題なく適用できる。すなわち、図10(c)の場合、データとしては、3つの音符 $n1^*$, $n2^*$, $n3^*$ が存在しているので、たとえば、音符 $n1^*$ のデータのみを残し、音符 $n2^*$, $n3^*$ のデータを削除する処理を行えばよい。

【0118】

図11は、このような重複回避処理を伴う音高シフトの別な例を示す図である。ここでは、図11(a)に示すとおり、5つの音符 $n1 \sim n5$ からなる和音について、本発明に係る方法で1オクターブ上げる音高シフトを行った例が示されている。この例では、オフセット値 $\alpha = 12$ であるが、各音符のノートナンバー n に応じて定まる重み関数値 $W(n)$ は、ノートナンバー n の増加に応じて徐々に減少してゆくため、各音符のシフト量は図示のとおり、それぞれ12, 9, 7, 4, 2となっている。

【0119】

図11(b)は、これらのシフト量に従って、各音符 $n1 \sim n5$ をシフトすることによって得られる新たな音符 $n1^* \sim n5^*$ を示している（破線の音符は、図11(a)に示す音符 $n1 \sim n5$ である）。なお、図示の便宜上、シフト後の5つの音符を左右2列に分けて描いているが、これら5つの音符 $n1^* \sim n5^*$ は、本来は、時間軸上の同一位置に配置されるべき音符である。図示のとおり、これら5つの音符のうちの一部は、音高変更後と同じ音高を占めている。すなわち、音符 $n1^*$ と音符 $n2^*$ とは同じ音高を占め、音符 $n3^*$ と音符 $n4^*$ とは同じ音高を占めている。

【0120】

そこで、重複した複数の音符については、1つのみを残して他を削除する重複回避処理を行う。図11(c)は、このような重複回避処理後の結果を示している。この例では、括弧書きで示した音符 $n2^*$ と音符 $n4^*$ とが削除されており、結局、音高変更後の音は、3つの音符 $n1^*$, $n3^*$, $n5^*$ の和音として表されている。

【0121】

なお、このような重複回避処理を行う場合、一部の符号の削除により、和音全体の強度は低下してしまう。すなわち、音高シフトを行う前は、図11(a)に示すように、5つの音符 $n1 \sim n5$ の和音によって構成されていた音が、重複回避処理を伴う音高シフトを行った後は、図11(c)に示すように、3つの音符 $n1^*$, $n3^*$, $n5^*$ の和音によって構成される音になってしまう。そこで、音の強度を維持する上では、重複回避処理を行う際に、1つのみ残された符号についての強度を、削除された符号についての強度に応じて修正するのが好ましい。

【0122】

MIDI符号の場合、音の強度の情報は、「ベロシティー」というデータとして符号に含まれているので、上例の場合、残された音符 $n1^*$ についてのベロシティーに、削除された音符 $n2^*$ についてのベロシティーを加える修正処理を施し、同様に、残された音符 $n3^*$ についてのベロシティーに、削除された音符 $n4^*$ についてのベロシティーを加える修正処理を施せばよい。このように、削除された符号についての強度に応じて強度修正を行えば、音高の変更処理後も適切な強度バランスをもった符号化データが得られる。

【0123】

<<< § 5. 本発明に係る符号化音声データの音高変換装置 >>>

図12は、本発明の基本的実施形態に係る符号化音声データの音高変換装置100の構成を示すブロック図である。この装置は、特定周波数の音が特定時間だけ持続することを示す符号を時間軸上に並べることによって構成され人間の声を表現する符号化音声データD（たとえば、MIDIデータ）を、異なる音高をもった別な音声データD*に変換する

10

20

30

40

50

符号化音声データの音高変換装置であり、その基本原理は、既に § 3, § 4 で述べたとおりである。

【0124】

図示のとおり、この音高変換装置 100 は、変換対象データ入力ユニット 110、オフセット値入力ユニット 120、音高変換処理ユニット 130、重み関数格納ユニット 140 によって構成される。変換対象データ入力ユニット 110 は、変換対象となる符号化音声データ D を入力する構成要素であり、オフセット値入力ユニット 120 は音高に関するオフセット値 α を入力する構成要素である。また、音高変換処理ユニット 130 は、変換対象となる符号化音声データ D に対して、オフセット値 α に基づく音高の変更処理を行い、変更後の符号化音声データ D* を出力する構成要素であり、重み関数格納ユニット 140 は、周波数 f について定義された所定の重み関数 $W(f)$ を格納する構成要素である。

10

【0125】

音高変換処理ユニット 130 は、§ 3 で述べた基本原理に基づいて、音高変換処理（音高シフト処理）を行う。具体的には、重み関数格納ユニット 140 に格納されている重み関数 $W(f)$ を用いて、変換対象となる符号化音声データ D に含まれている個々の符号について、当該符号が示す周波数 f を $\alpha \cdot W(f)$ に応じた値だけ増減することにより新たな周波数 f' を求め、当該符号を、それが示す周波数 f を f' に変更した新たな符号に置き換える処理を行う。

【0126】

ここで、重み関数格納ユニット 140 に格納されている重み関数 $W(f)$ は、図 7(a) に例示するように、周波数 f 軸上の所定区間において $W(f)$ が周波数 f の増加に従って単調減少する関数である。より具体的には、§ 3 で述べたとおり、周波数 f 軸上の第 1 設定値 f_a および第 2 設定値 f_b ($f_b > f_a$) について、 $f \leq f_a$ の区間は、 $W(f) = 1$ 、 $f_a < f < f_b$ の区間は、 $1 > W(f) > 0$ （但し、 $W(f)$ は f の増加に従って単調減少、たとえば、周波数 f の対数値に対して反比例）、 $f \geq f_b$ の区間は、 $W(f) = 0$ となる重み関数 $W(f)$ が用いられている。

20

【0127】

第 1 設定値 f_a は、平均的な音声についての基本周波数成分が含まれると予想される周波数領域の上限に設定し、第 2 設定値 f_b は、第 1 設定値 f_a に対して十分に高い任意の値に設定するのが好ましい。実際には、基本周波数成分が含まれると予想される周波数領域は男女でかなりの差があるので、実用上は、重み関数格納ユニット 140 には、男性用重み関数 $W\text{-male}(f)$ と女性用重み関数 $W\text{-female}(f)$ とを用意しておき、変換対象となる符号化音声データ D の声主の性別により、いずれか一方の重み関数を選択して利用するのが好ましい。

30

【0128】

すなわち、音高変換処理ユニット 130 は、変換対象となる符号化音声データ D によって表現される声が男性の声か女性の声かを示すオペレータからの指示に基づいて、男性の声の場合には男性用重み関数 $W\text{-male}(f)$ を用いた変更処理を行い、女性の声の場合には女性用重み関数 $W\text{-female}(f)$ を用いた変更処理を行うことになる。

【0129】

本願発明者が行った実験によれば、男性用重み関数 $W\text{-male}(f)$ については、第 1 設定値 f_a を 100 Hz ~ 200 Hz の範囲内に設定し、第 2 設定値 f_b を 3 kHz ~ 6 kHz の範囲内に設定すれば、良好な結果が得られた。また、女性用重み関数 $W\text{-female}(f)$ については、第 1 設定値 f_a を 200 Hz ~ 400 Hz の範囲内に設定し、第 2 設定値 f_b を 4 kHz ~ 8 kHz の範囲内に設定すれば、良好な結果が得られた。

40

【0130】

もちろん、必要に応じて、子供用重み関数を用意したり、年齢別の重み関数を用意したり、言語ごとに異なる重み関数を用意してもよい。

【0131】

オフセット値入力ユニット 120 が入力するオフセット値 α は、符号化音声データ D に

50

対する音高の変更指示を示す値であり、音高をどれだけ上げるのか、もしくは下げるのか、を示すことができる値であれば、どのように定義された値であってもかまわないが、一般的には、符号をもった数とし、音高を高める場合は正、低める場合は負の値をとり、絶対値が大きいほど、音高の変更量が大きいことを示す値として定義するのが好ましい。

【0132】

重み関数 $W(f)$ の値域を「 $1 \geq W(f) \geq 0$ 」とし、オフセット値 α を、上例のように定義した場合、音高変換処理ユニット 130 は、所定の係数 k ($k > 1$) を用いた式 $f' = f \cdot k^{\alpha \cdot W(f)}$ により新たな周波数 f' を求めることができる。 α が正負いずれの場合も、 $W(f) = 0$ の場合は、 $f' = f$ となり、周波数の変更は生じないので、第 2 設定値 f_b 以上の周波数をもつ符号に対しては、音高のシフト処理は行われなくなる。また、 $W(f) = 1$ の場合は、 $f' = f \cdot k^{\alpha}$ となり、オフセット値 α に対応するシフト量をもった周波数変更が行われる ($\alpha = +1$ であれば、周波数は k 倍になり、 $\alpha = -1$ であれば、周波数は $1/k$ 倍になる)。また、 $1 > W(f) > 0$ の場合には、周波数 f が大きいほど、周波数変更のシフト量は小さくなる。

10

【0133】

なお、§ 4 で述べたとおり、MIDI データのように、周波数 f をその対数値に対応するノートナンバー n で表す符号データを用いる場合、すなわち、変換対象データ入力ユニット 110 が入力した符号化音声データ D が、周波数 f をノートナンバー n によって示す符号を含むデータである場合、音高のシフト処理は、実際には、ノートナンバーの増減によって行うことになる。この場合、重み関数格納ユニット 140 には、図 7 (b) に例示するように、ノートナンバー n について定義された重み関数 $W(n)$ を格納しておけばよい。図 7 (b) に示す例は、周波数 f 軸上の第 1 設定値 f_a に対応するノートナンバー n_a および第 2 設定値 f_b に対応するノートナンバー n_b について、 $n_a < n < n_b$ の区間は、 $W(n)$ の値がノートナンバー n に反比例する値となる重み関数 $W(n)$ の例である。このような重み関数 $W(n)$ を用いると、前述したとおり、演算負担は大幅に軽減される。

20

【0134】

一方、オフセット値入力ユニット 120 によって入力されるオフセット値 α としては、音高を高める場合は正、低める場合は負の値をとるノートナンバーの差を用いるようにすればよい。たとえば、音高を 1 オクターブ上げる場合は、 $\alpha = +12$ 、音高を 1 オクターブ下げる場合は、 $\alpha = -12$ というオフセット値を与えればよい。

30

【0135】

また、音高変換処理ユニット 130 は、変換対象となる符号化音声データ D に含まれる個々の符号について、当該符号が示すノートナンバー n を用いた「 $n' = n + \alpha \cdot W(n)$ 」なる演算式により新たなノートナンバー n' を求め、当該符号を、それが示すノートナンバー n を n' に変更した新たな符号に置き換える処理を行うことになる。たとえば、 $\alpha = +12$ であった場合 (1 オクターブ上げる指示)、新たなノートナンバー n' は、「 $n' = n + 12 \cdot W(n)$ 」で与えられ、 $\alpha = -12$ であった場合 (1 オクターブ下げる指示)、「 $n' = n - 12 \cdot W(n)$ 」で与えられる。

【0136】

§ 1 で述べた方法により人間の音声を符号化した MIDI データは、音の持続時間が時間軸上で同一期間を占め、互いに異なる周波数を示す複数の符号を含む符号化音声データであり、いわば複数の音符からなる和音として音声を表現したデータである。したがって、変換対象データ入力ユニット 110 が、符号化音声データ D として MIDI 規格のデータを入力した場合、音高変換処理ユニット 130 が行う処理は、和音を構成する複数の符号のそれぞれについて、新たな符号への置換を行う処理ということになる。

40

【0137】

既に述べたとおり、MIDI 規格では、同じ音高をもつ複数の符号が重複して存在することは許されないため、音高変換処理ユニット 130 は、音の持続時間が時間軸上で同一期間を占める複数の符号についてそれぞれ新たな符号への置換処理を行う際に、同一のノートナンバーを示す新たな符号が複数 m 個生じた場合には、当該複数 m 個の符号のうち 1

50

つのみを残し、その余の $(m-1)$ 個を削除する重複回避処理を行う機能を有している。また、MIDIデータには、音の強度の情報（ベロシティ）が含まれているので、音高変換処理ユニット130は、重複回避処理を行う際に、1つのみ残された符号についての強度を、削除された符号についての強度に応じて修正する処理を行うようにするのが好ましい。

【0138】

なお、変換対象データ入力ユニット110によって入力された符号化音声データDがMIDIデータであった場合、音高変換処理ユニット130から出力される符号化音声データD*もMIDI規格のデータとすれば、この音高変換装置100は、符号化音声データの規格は維持したまま、音高のみを変換する処理を行うことができる。すなわち、ユーザは、所望のMIDIデータDをこの音高変換装置100に与え、所望のオフセット値 α を指定する操作を行えば、音高変換処理が施されたMIDIデータD*を得ることができる。

10

【0139】

こうして得られたMIDIデータD*は、もちろん、電子楽器による演奏の対象とすることができ、音高変換後の人間の音声を再生することができる。また、音高変換処理ユニット130に、符号化音声データD*を、五線譜上に音符を配置した楽譜のデータとして出力する機能をもたせておけば、音高変換後の人間の音声を表現した楽譜を得ることができ、実際の楽器を用いて、この楽譜上の音声情報を演奏して再生することもできる。

【0140】

20

<<< § 6. 本発明に係る音声変換装置 >>>

ここでは、§5で述べた本発明に係る音高変換装置100を利用して、音声変換装置を構成した例を、図13のブロック図を参照しながら説明する。図13に示す音声変換装置は、図12に示す音高変換装置100に、音声符号化装置200を更に付加することによって構成される装置である。

【0141】

図示のとおり、音声符号化装置200は、音声信号入力ユニット210と符号化ユニット220とによって構成される。ここで、音声信号入力ユニット210は、人間の声を含む音声信号Sをアナログ信号もしくはデジタル信号として入力する構成要素である。図には、オーディオCD再生器や一般のパソコンなどのオーディオ機器10から、デジタル音声信号Sを取り込む例と、マイク20で集音してアンプ装置30で増幅したアナログ音声信号Sを取り込む例とが示されている。

30

【0142】

一方、符号化ユニット220は、こうして取り込んだ音声信号Sを、特定周波数の音が特定時間だけ持続することを示す符号を時間軸上に並べることによって構成された符号化音声データDに変換する処理を行う構成要素である。具体的には、§1で述べた方法によって、人間の音声をMIDI符号化する装置（たとえば、前掲の特許文献1～3に開示されている装置）によって、符号化ユニット220を構成することができる。

【0143】

こうして、符号化ユニット220から出力された符号化音声データD（たとえば、MIDIデータ）は、音高変換装置100に与えられる。この音高変換装置100は、図12に示す構成要素であり、その機能は既に§5で述べたとおりである。すなわち、所望のオフセット値 α を与えることにより、入力された符号化音声データDに対して音高変更処理を施し、符号化音声データD*を出力する機能を有している。

40

【0144】

結局、この音声変換装置を利用すると、オーディオ機器10によって再生した歌手の歌声を、デジタル音声信号Sとして与えることにより、当該歌声を表現したMIDI形式等の符号化音声データD*を得ることができ、しかもその音高を、与えるオフセット値 α によって任意に変化させることができる。あるいは、マイク20でリアルタイムに集音した人間の音声を、その場でMIDI形式等の符号化音声データD*に変換することができ、

50

しかもその音高を、与えるオフセット値 α によって任意に変化させることができる。要するに、入力した人間の音声の音高を自由に変更した上で、符号化音声データ D^* として出力することができる。

【0145】

もちろん、こうして出力された符号化音声データ D^* は、MIDI音源等を利用して即座に再生することが可能であり、楽譜として出力すれば、楽器を用いて演奏することが可能である。

【0146】

<<< § 7. 本発明に係る音声合成装置 >>>

ここでは、§ 5で述べた本発明に係る音高変換装置100を利用して、音声合成装置を構成した例を、図14のブロック図を参照しながら説明する。図14に示す音声合成装置は、図12に示す音高変換装置100に、テキスト符号化装置300を更に付加することによって構成される装置である。このテキスト符号化装置300の詳細な構成や動作は、前掲の特許文献4に記載されているが、ここではその原理だけを簡単に述べておく。

【0147】

図示のとおり、テキスト符号化装置300は、符号データベースユニット310と、テキストデータ入力ユニット320と、符号合成ユニット330と、によって構成される。ここで、符号データベースユニット310は、所定の言語による単語を構成する個々の音節にそれぞれ対応する符号群（特定周波数の音が特定時間だけ持続することを示す符号の集合体）を格納した構成要素であり、テキストデータ入力ユニット320は、所定の言語による単語を文字列として表現したテキストデータ T を入力する構成要素であり、符号合成ユニット330は、符号データベースユニット310を参照して、入力されたテキストデータ T の読みを構成する個々の音節にそれぞれ対応する符号群を抽出し、これらを時間軸上に並べることによって、テキストデータ T の読みに対応する人間の声を表現した符号化音声データ D を合成する構成要素である。

【0148】

こうして、符号合成ユニット330で合成された符号化音声データ D （たとえば、MIDIデータ）は、音高変換装置100に与えられる。この音高変換装置100は、図12に示す構成要素であり、その機能は既に§ 5で述べたとおりである。すなわち、所望のオフセット値 α を与えることにより、入力された符号化音声データ D に対して音高変更処理を施し、符号化音声データ D^* を出力する機能を有している。

【0149】

図13に示す音声変換装置が、音声信号 S に基づいて符号化音声データ D^* を生成する機能を有する装置であるのに対して、この図14に示す音声合成装置は、テキストデータ T に基づいて符号化音声データ D^* を生成する機能を有する装置ということができる。そのため、符号データベースユニット310には、所定の言語による単語を構成する個々の音節にそれぞれ対応する符号群が格納されている。特に、ここに示す実施形態の場合、日本語による単語を構成する個々の音節にそれぞれ対応する符号群として、子音を構成する子音音素に対応する符号群（子音音素用 DB ）と、母音を構成する母音音素に対応する符号群（母音音素用 DB ）と、が格納されている。

【0150】

符号合成ユニット330は、テキストデータ T の読みを構成する個々の音節を子音音素と母音音素とに分解し、符号データベースユニット310を参照して、個々の音素ごとにそれぞれ対応する符号群を抽出し、これらを時間軸上に並べることによって符号化音声データ D を合成する処理を行う。

【0151】

図15は、符号合成ユニット330内に用意されている変換テーブルを示す図である。この変換テーブルは、日本語の音節を構成する71個のカタカナを、それぞれ子音音素と母音音素とに分解するためのテーブルであり、各カタカナは、テーブルの1行目に記載された子音音素（ K , S , T , ... 等の合計14種類）と、テーブルの1列目に記載された

10

20

30

40

50

母音音素（A，I，U，E，O，nの合計6種類）とに分解される。なお、カタカナの「ア，イ，ウ，エ，オ，ン」は、母音音素のみから構成される。また、「キャ」，「シュ」などの拗音を含むカナや、「トッ」のような促音を含むカナは、これらの組み合わせとして取り扱われる。

【0152】

たとえば、テキストデータ入力ユニット320によって、「コンニチワ」なる5文字のカタカナから構成されるテキストデータTが入力された場合、符号合成ユニット330は、図15に示す変換テーブルを利用して、この「コンニチワ」の5音節を、それぞれ子音音素と母音音素とに分解して、図16に示すように、「KOnNITiWA」なる9音素を生成する。

10

【0153】

一方、符号データベースユニット310内には、図17(a)～(d)に示すようなMIDI符号データベース（男性）と、図18(a)～(d)に示すようなMIDI符号データベース（女性）とが用意されている。これらのデータベースは、6種類の母音音素「A，I，U，E，O，n」と14種類の子音音素「K，S，T，N，H，M，R，G，Z，D，B，P，Y，W」について、それぞれ所定の符号群を対応づけたものである。すなわち、各音素は、それぞれ8個の音符からなる符号群に対応づけられている。たとえば、図17(a)の「A」の欄には、「F5，E5，D#5，A4，G#4，C2，B1，A#1」なる8個の記号が収容されているが、これら8個の記号は、五線譜上の音階を示す記号である。したがって、この例の場合、1つの音素は、五線譜上において、8個の音符の和音として表現されることになる。なお、この図17や図18に示すようなデータベースは、アナウンサーや声優などに依頼して録音したサンプルについて、§1で述べた符号化方法を適用することにより準備することができる。

20

【0154】

図19および図20は、それぞれ図17に示す男性用のMIDI符号データベースおよび図18に示す女性用のMIDI符号データベースを用いて、各母音音素と各子音音素とに対応する符号を、音符として表示した例を示す楽譜である。実際には、符号合成ユニット330は、MIDIデータから構成される符号化音声データDを生成する処理を実行する。なお、図17および図18に示すデータベースには、強度や時間の情報は含まれていないので、MIDIデータを生成する際に必要なベロシティやデルタタイムのデータは、たとえば、予め設定した固定値を用いるようにすればよい。

30

【0155】

結局、符号合成ユニット330は、「コンニチワ」の5音節を「KOnNITiWA」なる9音素に分解した後、符号データベースユニット310を参照して、個々の音素ごとにそれぞれ対応する符号群を抽出し、これらを時間軸上に並べることによって符号化音声データDを合成する処理を行うことができる。ここに示す例では、符号データベースユニット310には、男性用のデータベース（図17）と女性用のデータベース（図18）とが用意されているので、オペレータにいずれか一方を選択させ、選択されたデータベースを利用して、符号群の抽出を行うようにする。もちろん、同じ男性用のデータベースであっても、低音声の男性用、高音声の男性用のように別々のデータベースを用意することもできる。

40

【0156】

こうして、上例の場合、テキスト符号化装置300によって、「コンニチワ」なるテキストデータTに基づいて、この5音節を男性もしくは女性が発音した場合に対応する符号化音声データD（MIDIデータ）が出力される。この符号化音声データDは、音高変換装置100に与えられるので、所望のオフセット値 α を与えることにより、音高変更処理を施した符号化音声データD*を得ることができる。

【0157】

この音声合成装置を利用すると、カナ文字からなる任意のテキストデータTに基づいて、当該カナ文字の読みを表現したMIDI形式等の符号化音声データD*を得ることがで

50

き、しかもその音高を、与えるオフセット値 α によって任意に変化させることができる。もちろん、こうして出力された符号化音声データ D^* は、MIDI音源等を利用して再生することが可能であり、楽譜として出力すれば、楽器を用いて演奏することが可能である。また、符号データベースとして外国語の音節に対応するデータベースを用意しておけば、外国語のテキストに基づいて符号化音声データを得ることも可能である。

【0158】

<<< § 8. いくつかの変形例 >>>

ここでは、これまで述べてきた様々な実施形態についての変形例を述べる。

【0159】

< 8-1. 音声合成装置の第1の変形例 >

10

図21は、図14に示す音声合成装置の第1の変形例を示すブロック図である。図14に示す音声合成装置では、音高を変更するためには、音高変換装置100に対してオフセット値 α を指定する必要があるが、この第1の変形例に係る音声合成装置では、予めテキストデータ内にオフセット値 α を埋め込んでおくことにより、テキスト符号化装置300Aから音高変換装置100に対して、自動的に、かつ、音素ごとにオフセット値 α を指示することができる。

【0160】

図21に示すとおり、テキスト符号化装置300Aは、符号データベースユニット310と、テキストデータ入力ユニット320と、符号合成ユニット340と、によって構成されている。ここで、符号データベースユニット310とテキストデータ入力ユニット320は、図14に示すテキスト符号化装置300に用いられているものと同じものでかまわない。ただ、図14に示すテキスト符号化装置300の場合、入力対象となるテキストデータTとして、図22(a)に例示する「コンニチワ」のように、カナ文字からなる文字列を用いていたが、この図21に示すテキスト符号化装置300Aの場合、所定のオフセット値 α が埋め込まれたテキストデータT(α)を用いることになる。

20

【0161】

図22(b)は、オフセット値 α が埋め込まれたテキストデータT(α)の一例を示す図である。例示したテキストデータT(α)は、「-3コ+3ン+6ニ+3チ-3ワ」のように、「半角の+もしくは-」、「半角数字」、「全角カナ文字」によって構成される。ここで、「全角カナ文字」は、図22(a)に示す「コンニチワ」なる文字列であり、本来のテキストデータを構成する文字列である。一方、「半角の+もしくは-」および「半角数字」は、個々の「全角カナ文字」の前に挿入された記号であり、後続するカナ文字に対するオフセット値 α を示すものである。たとえば、図22(b)に示す例において、「-3」は、後続するカナ文字「コ」についてのオフセット値 $\alpha = -3$ を示しており、「+3」は、後続するカナ文字「ン」についてのオフセット値 $\alpha = +3$ を示している。

30

【0162】

結局、図22(b)に示すテキストデータT(α)は、個々の「全角カナ文字」の前に、当該カナ文字のオフセット値 α を示すための記号として、「半角の+もしくは-」および「半角数字」を挿入する、という書式をもったテキストデータということになる。テキストデータ入力ユニット320は、このような音節ごとのオフセット値 α を含むテキストデータT(α)を取り込み、そのまま符号合成ユニット340へと与える。

40

【0163】

一方、符号合成ユニット340は、上記書式に基づいて、与えられたテキストデータT(α)を、オフセット値 α とカナ文字とに分離し、カナ文字については、個々の音節を子音音素と母音音素とに分解し、符号データベースユニット310を参照して、個々の音素ごとにそれぞれ対応する符号群を抽出し、これらを時間軸上に並べることによって符号化音声データDを合成する(この処理は、図14に示す音声合成装置300内の符号合成ユニット330が行う処理と全く同様である)。そして、こうして合成した符号化音声データDとともに、個々の音節ごとのオフセット値 α を音高変換装置100に対して出力する。

50

【0164】

たとえば、図22(b)に示すような「-3コ+3ン+6ニ+3チ-3ワ」というテキストデータT(a)が与えられた場合、符号合成ユニット340は、これを「コ、ン、ニ、チ、ワ」なる5つの音節と、各音節に対応するオフセット値「-3, +3, +6, +3, -3」とに分離し、「コ、ン、ニ、チ、ワ」なる音節を更に、「K, O, n, N, I, T, I, W, A」なる9つの音素に分解し、個々の音素に対して、図22(c)に示すように、元の音節に対応するオフセット値 α を割り当てる。そして、「K, O, n, N, I, T, I, W, A」なる9つの音素について、符号データベースユニット310を参照して、対応する符号群を抽出し、符号化音声データ(MIDIデータ)を生成する。ここでは、音素xについて生成した符号化音声データをD(x)と表すことにする。最後に、これら各符号化音声データをD(x)に、割り当てられたオフセット値 α を示す情報を付加し、音高変換装置100へと出力する。

10

【0165】

図22(d)は、上例の場合に、符号合成ユニット340から出力されるデータを示す図である。たとえば、1行目の符号化音声データD(K)は、子音音素「K」に対応するデータであり(たとえば、図17(b)の[K]欄に記載された8個の音符に対応するMIDIデータ)、それに後続する「 $\alpha = -3$ 」は、当該データD(K)に割り当てられたオフセット値 α が-3であることを示すデータである。図示のとおり、「コンニチワ」なる単語を構成する9つの音素のすべてについて、それぞれ符号化音声データとオフセット値 α とが与えられている。

20

【0166】

この図22(d)に示すような「個々の音節(音素)ごとのオフセット値 α を伴う符号化音声データD」が与えられた音高変換装置100は、個々の音節(音素)ごとに、割り当てられているオフセット値 α を用いた音高シフト処理を実行する。たとえば、図22(d)の1行目に示す符号化音声データD(K)については、オフセット値 $\alpha = -3$ に応じた音高シフト処理が実行される。データD(K)を構成する符号が、たとえば、図17(b)の[K]欄に記載された8個の音符に対応するMIDIデータであった場合、この8個の音符のノートナンバーに対して、オフセット値 α に応じた増減処理(元のノートナンバーをnとしたときに、「 $n' = n + \alpha \cdot W(n)$ 」で与えられる新たなノートナンバーn'に変更する処理)が行われる。

30

【0167】

このように、音高変換装置100は、符号合成ユニット340から与えられたオフセット値 α を用いて、個々の音節ごとに音高の変更処理を行うことができるので、同じ「コンニチワ」という単語であっても、音節ごとのオフセット値 α を変化させることによって、自然な抑揚をもった標準語のアクセントによる「コンニチワ」や、関西弁のアクセントによる「コンニチワ」などのバリエーションをもった符号化音声データD*を生成することができる。もちろん、通常は耳にしないような奇妙なアクセントをもった「コンニチワ」を生成することも可能なので、エンターテインメント分野での利用に適している。

【0168】

<8-2. 音声合成装置の第2の変形例>

40

図23は、図14に示す音声合成装置の第2の変形例を示すブロック図であり、テキスト符号化装置300の代わりに、テキスト符号化装置300Bが用いられている。テキスト符号化装置300Bは、図14に示すテキスト符号化装置300における符号合成ユニット330を、付加機能をもった符号合成ユニット350に置き換え、新たに、音高辞書ユニット360を追加したものである。符号データベースユニット310とテキストデータ入力ユニット320は、図14に示すテキスト符号化装置300に用いられているものと同じものでかまわない。

【0169】

図14に示す音声合成装置では、音高変換装置100に対して、外部からオフセット値 α を指示する必要があった。一方、図21に示す第1の変形例に係る音声合成装置では、

50

オフセット値 α が埋め込まれたテキストデータ $T(\alpha)$ を与える必要があった。これに対して、図 2 3 に示す第 2 の変形例に係る音声合成装置では、オフセット値 α が自動的に付与されるので、外部からオフセット値 α を指定する必要がないというメリットが得られる。

【0170】

オフセット値 α の自動付与に貢献するのは、音高辞書ユニット 3 6 0 である。この音高辞書ユニット 3 6 0 は、多数の単語について、それぞれ当該単語を構成する各音節に与えるオフセット値 α を格納した辞書である。図 2 4 は、この音高辞書ユニット 3 6 0 の内容を例示する表である。図示の例では、左欄の各単語について、当該単語を構成する各音節に適したオフセット値 α が右欄に示されている。たとえば、「オハヨウ」という単語について、右欄には「 $-1, +6, +4, -3$ 」というオフセット値 α が記載されているが、これは、「オハヨウ」という単語を構成する 4 音節のそれぞれのオフセット値を示すものであり、音節「オ」については $\alpha = -1$ 、音節「ハ」については $\alpha = +6$ 、音節「ヨ」については $\alpha = +4$ 、音節「ウ」については $\alpha = -3$ であることを示している。

【0171】

この図 2 4 には、わずか 3 単語についてのオフセット値しか示されていないが、実際の音高辞書ユニットには、通常の辞書と同様に、多数の単語について、それぞれ当該単語を構成する各音節に適したオフセット値 α が格納されている。このような音高辞書ユニット 3 6 0 を用意するには、アナウンサーなどに標準語で各単語を発音してもらい、当該単語に含まれる各音節の音高（たとえば、スペクトルでのピーク周波数）を測定し、規準音高（予め設定した特定の周波数、たとえば、多数の単語の各音節についての平均音高）に対する音程差を、各音節のオフセット値 α とすればよい。もちろん、標準語辞書の他に、関西弁辞書、東北弁辞書などを用意することも可能である。

【0172】

さて、この変形例 2 に係る装置の場合、テキストデータ T としては、カナ文字からなる通常のテキストを与えればよく、オフセット値 α を含むテキストデータ $T(\alpha)$ を与える必要はない。たとえば、「コンニチワ」なる文字列からなるテキストデータ T をテキストデータ入力ユニット 3 2 0 に与えたとすると、当該テキストデータ T は、そのまま符号合成ユニット 3 5 0 へ伝えられる。符号合成ユニット 3 5 0 は、個々の音節を子音音素と母音音素とに分解し、符号データベースユニット 3 1 0 を参照して、個々の音素ごとにそれぞれ対応する符号群を抽出し、これらを時間軸上に並べることによって符号化音声データ D を合成する（この処理は、図 1 4 に示す音声合成装置 3 0 0 内の符号合成ユニット 3 3 0 が行う処理と全く同様である）。

【0173】

ただ、この符号合成ユニット 3 5 0 には、オフセット値 α の自動付与機能が備わっている。すなわち、符号合成ユニット 3 5 0 は、音高辞書ユニット 3 6 0 を参照することにより、単語「コンニチワ」についての個々の音節ごとのオフセット値 α を認識し、符号化音声データ D とともに、このオフセット値 α を音高変換装置 1 0 0 に与える処理を実行する。具体的には、単語「コンニチワ」については、図 2 4 に示す辞書を引くことにより、「 $-3, +3, +6, +3, -3$ 」なるオフセット値 α が得られるので、符号化音声データ D に、これらのオフセット値 α を付加したデータが、音高変換装置 1 0 0 に対して出力される。

【0174】

こうして符号合成ユニット 3 5 0 から出力されるデータは、本質的には、図 2 2 (d) に示すデータと全く同じになる。符号合成ユニット 3 5 0 は、辞書を引くことにより、図 2 2 (c) に示すように、個々の音素ごとに割り当てるべきオフセット値 α を認識することができるので、図 2 2 (d) に示すように、個々の音素ごとの符号化音声データ $D(x)$ と、これに割り当てられたオフセット値 α を、音高変換装置 1 0 0 に対して出力することができる。なお、テキストデータ T が文章からなる場合に、当該文章を構成するカナ文字をどのように区切って個々の単語として認識するか、という方法に関しては、既に形態素解析

などの言語解析手法として様々な手法が知られているので、ここでは詳しい説明は省略する。

【0175】

結局、変形例1に係る装置では、図22(b)に例示するように、オフセット値 α を含むテキストデータT(α)を与えることにより、個々の音節についてのオフセット値 α を指示していたのに対して、変形例2に係る装置では、予め音高辞書ユニット360を用意しておくことにより、オフセット値 α の指示を自動化したことになる。音高変換装置100は、変形例1に係る装置と同様に、符号合成ユニット350から与えられたオフセット値 α を用いて、個々の音節ごとに音高の変更処理を行うことができる。

【0176】

<8-3. 音声発声装置の付加>

図25は、本発明に係る音高変換装置100から出力された符号化音声データD*に基づいて、音声を発声させる音声発声装置400の基本構成を示すブロック図である。MIDIデータのような符号化音声データには、実際の音の波形に関する情報は含まれていないので、この符号化音声データを復号化して音を再生するためには、音の波形に関する情報をもった音源を用いる必要がある。

【0177】

図25に示す音声発声装置400は、図示のとおり、復号化ユニット410、音源ユニット420、発音ユニット430によって構成されている。ここで、音源ユニット420は、所定の楽器による様々な周波数の演奏音響波形をデジタルデータとして格納した構成要素である。復号化ユニット410は、与えられた符号化音声データD*を構成する個々の符号を、音源ユニット420に格納されている対応する演奏音響波形に置き換えることにより音声信号Sの復号化を行う。また、発音ユニット430は、こうして復号化された音声信号Sに基づいて音波を生成する機能を果たし、実際には、アンプ装置やスピーカなどによって構成される。

【0178】

これまで述べてきた音声変換装置(図13)や音声合成装置(図14, 図21, 図23)に、図25に示すような音声発声装置400を付加すれば、音高変更が行われた後の符号化音声データD*に基づいて、その場で実際の音を発声させることができるようになる。

【0179】

<8-4. MIDI以外の規格による符号化>

これまで述べた実施形態では、主として、本発明をMIDI規格の符号化音声データを用いて実施した例を述べてきた。しかしながら、本発明を実施する場合、必ずしもMIDI規格の符号化音声データを用いる必要はない。本発明は、特定周波数の音が特定時間だけ持続することを示す符号を時間軸上に並べることによって構成され人間の声を表現する符号化音声データであれば、どのような規格の符号化音声データを用いた場合でも、その作用効果が得られるものである。

【0180】

<8-5. コンピュータプログラムによる実現>

これまでの説明では、本発明に係る音高変換装置、音声変換装置、音声合成装置を、ブロック図として示し、個々のブロックで示す構成要素の集合体として示したが、実用上、これらの装置はコンピュータに専用のプログラムを組み込むことによって構成することができる。

【0181】

<<< § 9. 方法発明としての把握 >>>

これまで、本発明を装置として把握した説明を行ったが、本発明に係る基本思想は、方法発明としても捉えることができる。

【0182】

具体的には、本発明は、特定周波数の音が特定時間だけ持続することを示す符号を時間

10

20

30

40

50

軸上に並べることによって構成され人間の声を表現する符号化音声データについて、その抑揚を変換する符号化音声データの抑揚変換方法として捉えることができる。この方法は、コンピュータが実行する方法であり、コンピュータが、変換対象となる符号化音声データDを入力する変換対象データ入力段階と、コンピュータが、音高に関するオフセット値 α を入力するオフセット値入力段階と、コンピュータが、変換対象となる符号化音声データDに対して、オフセット値 α に基づく音高の変更処理を行い、変更後の符号化音声データD*を出力する音高変換処理段階と、を実行する。

【0183】

ここで、音高変換処理段階において、周波数 f 軸上の所定区間において $W(f)$ が周波数 f の増加に従って単調減少する所定の重み関数 $W(f)$ を利用して、変換対象となる符号化音声データDに含まれている個々の符号について、当該符号が示す周波数 f を $\alpha \cdot W(f)$ に応じた値だけ増減することにより新たな周波数 f' を求め、当該符号を、それが示す周波数 f を f' に変更した新たな符号に置き換える処理を行うことが特徴となる。

【0184】

また、本発明をMIDIデータに適用した場合は、人間の声を表現したMIDIデータを、抑揚の異なる別なMIDIデータに変換するMIDIデータの抑揚変換方法として捉えることができる。この方法は、コンピュータが実行する方法であり、コンピュータが、変換対象となるMIDIデータDを入力する変換対象データ入力段階と、コンピュータが、音高に関するオフセット値 α を入力するオフセット値入力段階と、コンピュータが、変換対象となるMIDIデータDに対して、オフセット値 α に基づく音高の変更処理を行い、変更後のMIDIデータD*を出力する音高変換処理段階と、を実行する。

【0185】

ここで、変換対象データ入力段階では、互いに異なるノートナンバーをもち、時間軸上の同一位置を占める複数のMIDI符号を含むMIDIデータDの入力を行い、音高変換処理段階では、ノートナンバー軸上の第1設定値 n_a および第2設定値 n_b ($n_b > n_a$) について、 $n \leq n_a$ の区間は、 $W(n) = 1$ 、 $n_a < n < n_b$ の区間は、 $1 > W(n) > 0$ (但し、 $W(n)$ は n の増加に従って単調減少)、 $n \geq n_b$ の区間は、 $W(n) = 0$ となる所定の重み関数 $W(n)$ を利用して、変換対象となるMIDIデータDに含まれている個々のMIDI符号について、当該MIDI符号が示すノートナンバー n に対して $\alpha \cdot W(n)$ に応じた値だけ加減算を行うことにより新たなノートナンバー n' を求め、当該MIDI符号を、それが示すノートナンバー n を n' に変更した新たなMIDI符号に置き換える処理を行うことにより、変更後のMIDIデータD*を生成することが特徴となる。

【0186】

このMIDIデータの抑揚変換方法では、音高変換処理段階で、新たなMIDI符号に置き換える処理を行う際に、時間軸上の同一位置を占め、同一のノートナンバーをもつ新たなMIDI符号が複数 m 個生じた場合には、当該複数 m 個のMIDI符号のうち1つのみを残し、その余の $(m-1)$ 個を削除する重複回避処理を行うのが好ましい。

【産業上の利用可能性】

【0187】

本発明に係る技術は、人間の声をMIDI規格などに基づいて符号化して取り扱うことを前提とする技術である。しかも、PCMのような符号化とは根本的に異なり、本発明で用いられる符号化音声データは、特定周波数の音が特定時間だけ持続することを示すいくつかの符号を時間軸上に並べることによって構成されるデータであるため、再生した場合、当然ながら、元の人間の声に対して、かなりかけ離れた音声を得られることになる。したがって、本発明は、人間の声を正確に記録し、再生するという用途に利用されるものではない。

【0188】

しかしながら、「人間の声を音符等の符号によって表現し、楽器によって演奏できるようにする」という着想は極めてユニークであり、エンターテインメントの分野では、様々

10

20

30

40

50

な利用形態が期待できる。しかも本発明では、オフセット値 α を指定することにより、音高を自由に調整することができるので、その用途は更に広がることになる。

【0189】

具体的には、本発明は、イベントの余興目的に行われる人間の音声再生を模倣した音楽作品の制作や作曲の支援に利用することができる。また、エンターテインメントの分野において、電子楽器を主体とした玩具（ロボットやぬいぐるみなども含む）、玩具型のアコースティック楽器（室内装飾用のミチチュアピアノなど）、オルゴールなどにも利用することができ、更に、携帯電話の着信メロディなどの音階再生媒体に対して音声合成機能を付加する用途にも適している。

【符号の説明】

【0190】

10：オーディオ機器
 20：マイク
 30：アンプ装置
 100：音高変換装置
 110：変換対象データ入力ユニット
 120：オフセット値入力ユニット
 130：音高変換処理ユニット
 140：重み関数格納ユニット
 200：音声符号化装置
 210：音声信号入力ユニット
 220：符号化ユニット
 300, 300A, 300B：テキスト符号化装置
 310：符号データベースユニット
 320：テキストデータ入力ユニット
 330：符号合成ユニット
 340：符号合成ユニット
 350：符号合成ユニット
 360：音高辞書ユニット
 400：音声発生装置
 410：復号化ユニット
 420：音源ユニット
 430：発音ユニット（アンプ装置・スピーカ）
 A：複素強度
 D：符号化音声データ
 D*：音高変更後の符号化音声データ
 D(x)：音素 x の符号化音声データ
 d1～d5：単位区間
 E：実効強度（エネルギー）
 e(i, j)：符号コード n(i, j) の実効強度
 F：サンプリング周波数
 F0：基本周波数
 F1：第1フォルマント周波数
 F2：第2フォルマント周波数
 F3：第3フォルマント周波数
 f：周波数
 f'：音高変更後の周波数
 f1～f3：個々の音符が示す周波数
 f1*～f3*：音高変更後の個々の音符が示す周波数
 fa：周波数の第1設定値

10

20

30

40

50

f_b : 周波数の第 2 設定値

n : ノートナンバー

$n_1 \sim n_5$: 音符／個々の音符が示すノートナンバー

$n_1^* \sim n_5^*$: 音高変更後の音符／個々の音符が示すノートナンバー

n_a : 周波数 f_a に対応するノートナンバー

n_b : 周波数 f_b に対応するノートナンバー

$n(i, j)$: 単位区間 d_i について抽出された第 j 番目の符号コード

S : 音声信号

T : テキストデータ

$T(\alpha)$: オフセット値 α を含むテキストデータ

$T_1 \sim T_3$: トラック

t : 時間

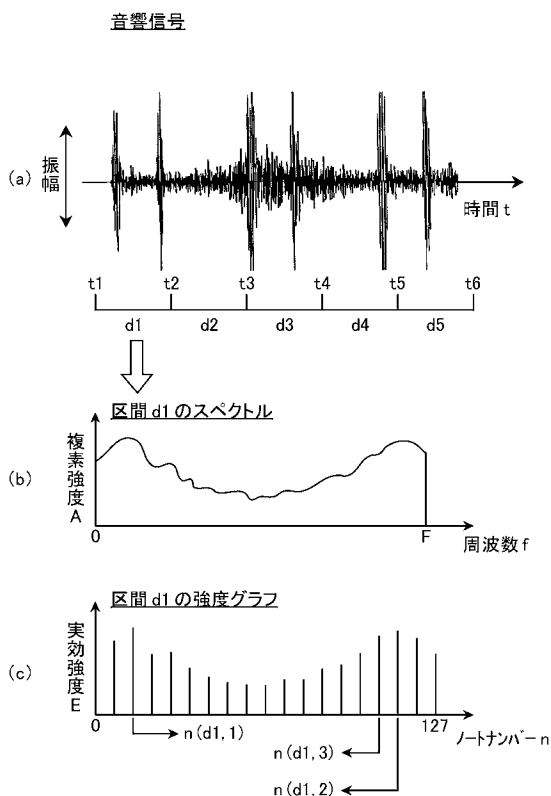
$t_1 \sim t_6$: 時刻

$W(f)$, $W(n)$: 重み関数

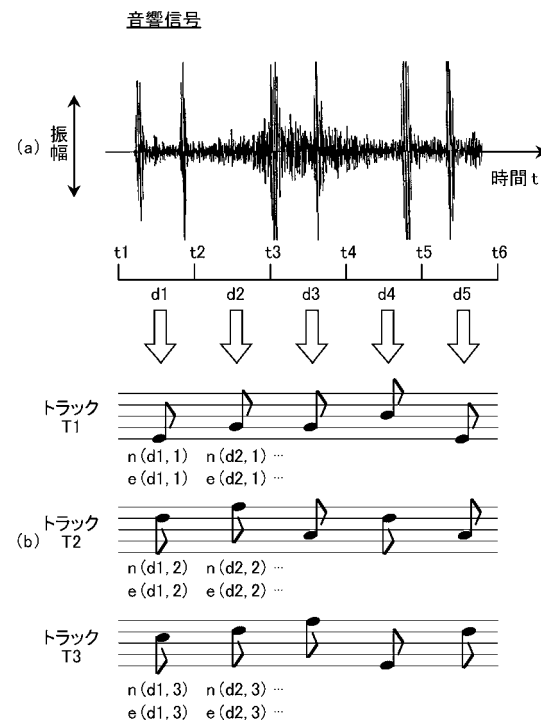
α : オフセット値

10

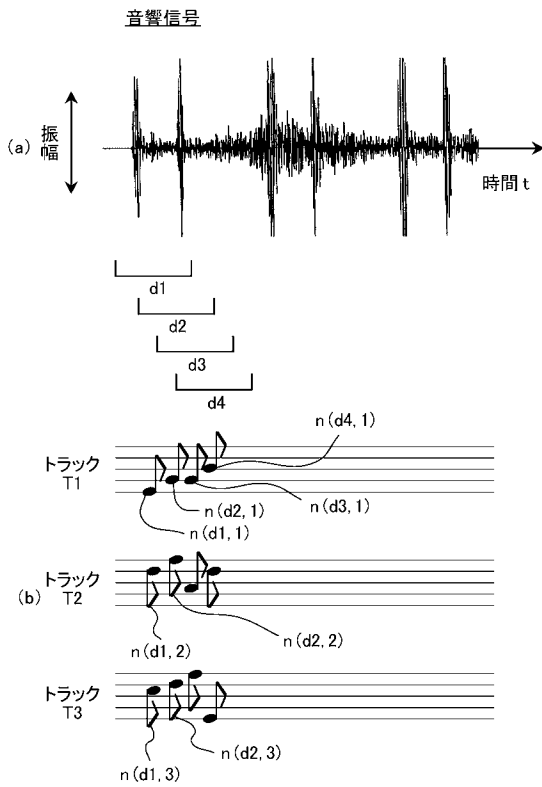
【図 1】



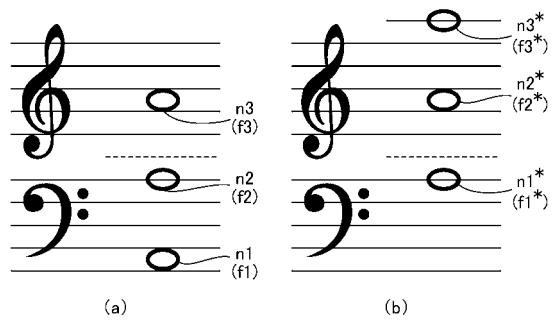
【図 2】



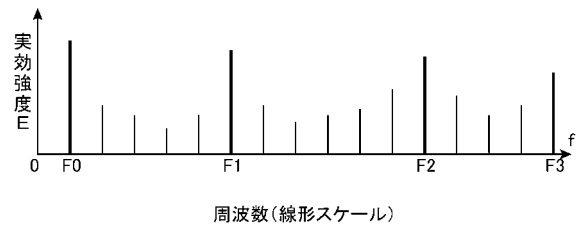
【図 3】



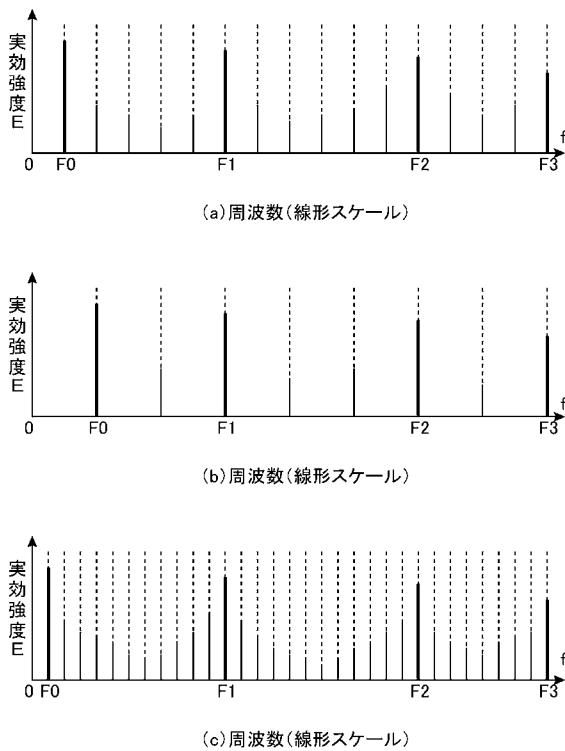
【図 4】



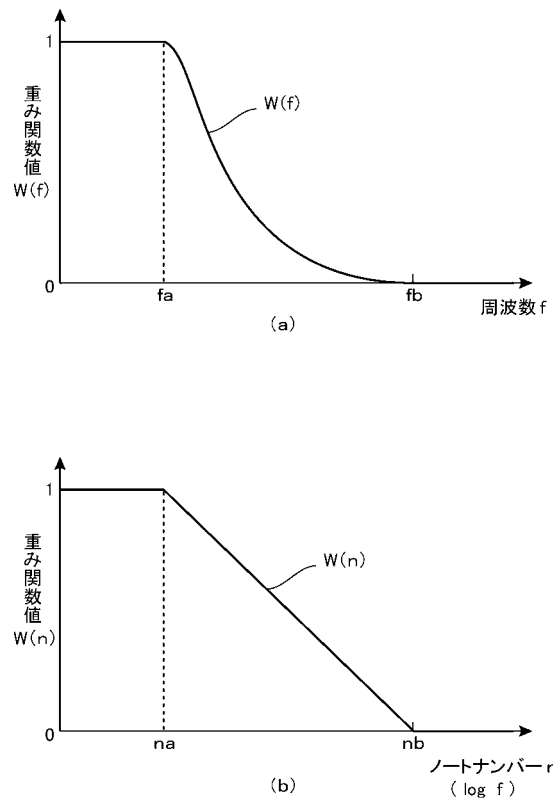
【図 5】



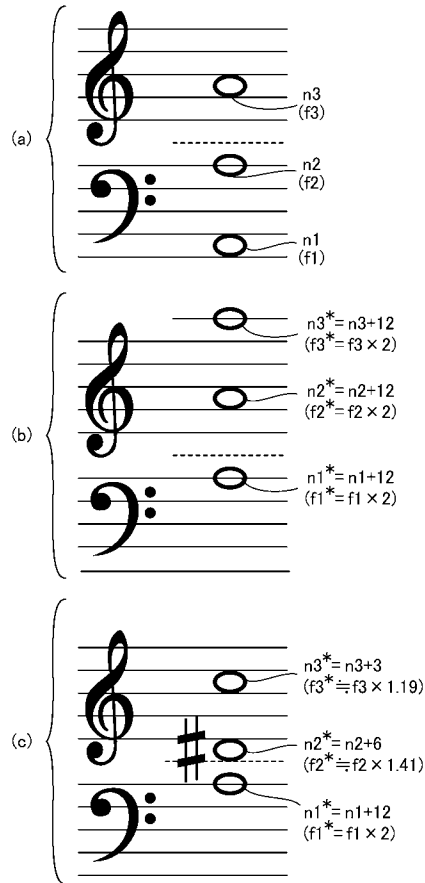
【図 6】



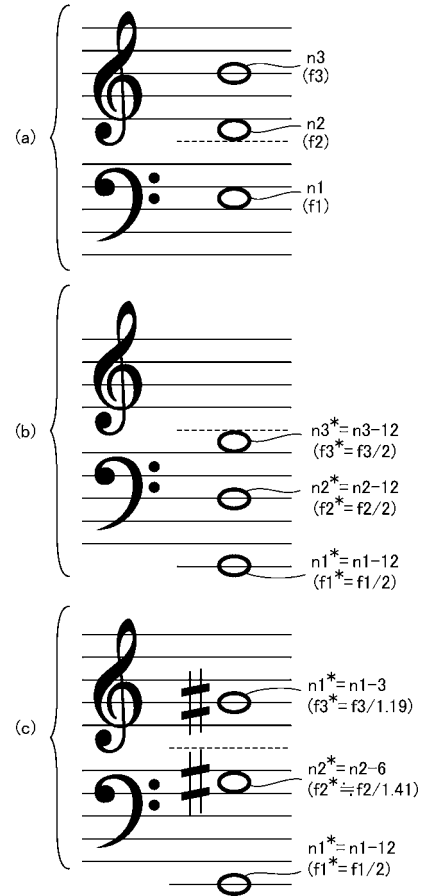
【図 7】



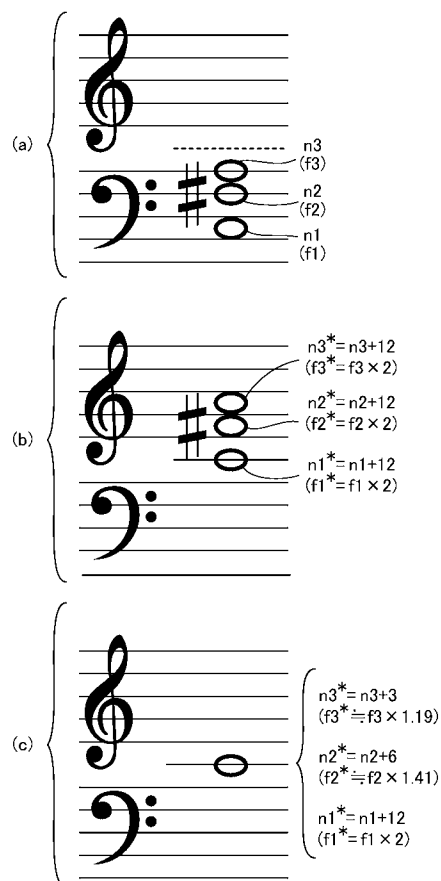
【図 8】



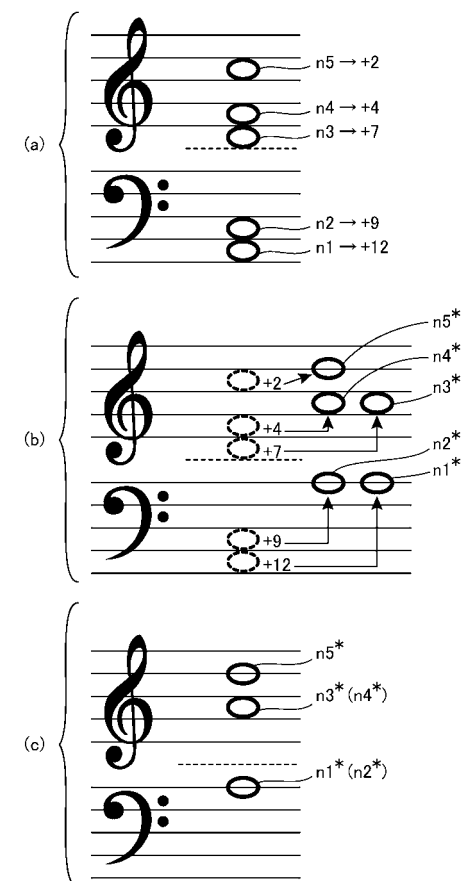
【図 9】



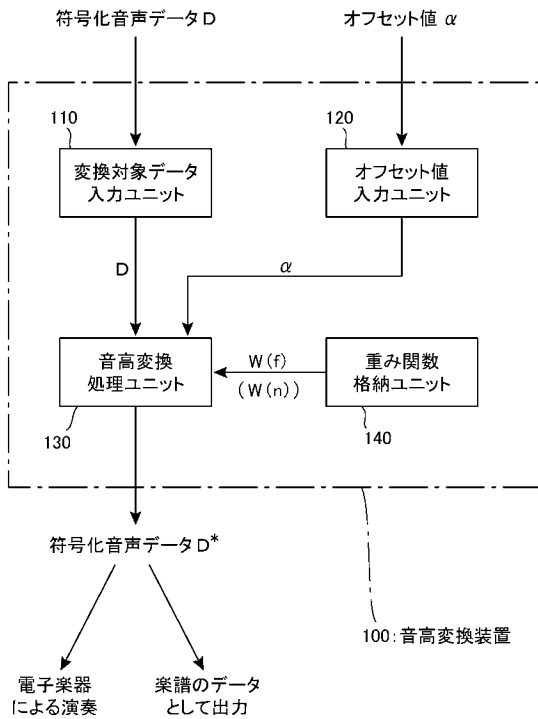
【図 10】



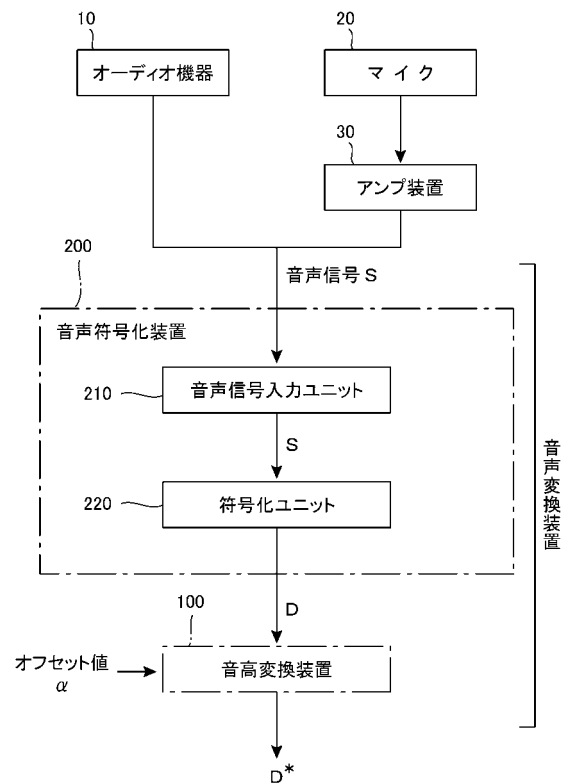
【図 11】



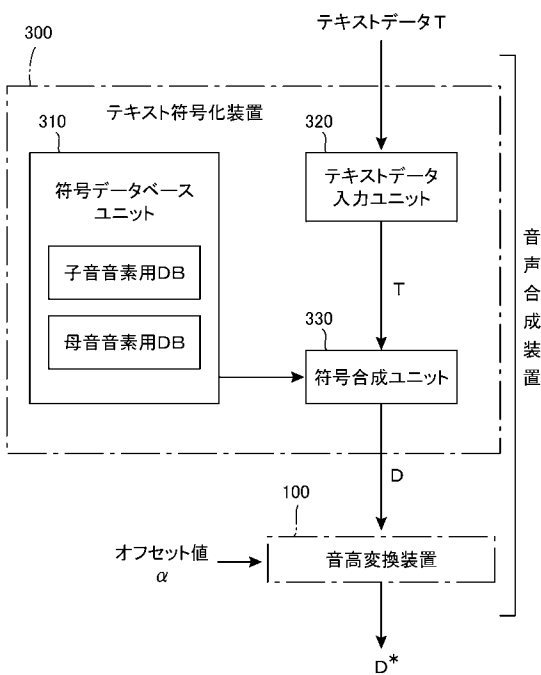
【図 1 2】



【図 1 3】



【図 1 4】



【図 1 5】

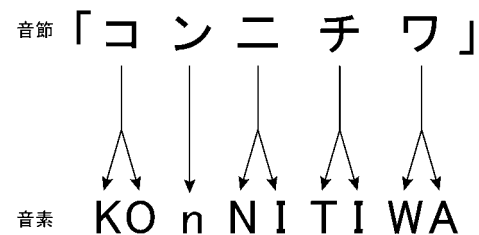
変換テーブル

	K	S	T	N	H	M	R	G	Z	D	B	P	Y	W
A	ア	カ	サ	タ	ナ	ハ	マ	ラ	ガ	ザ	バ	パ	ヤ	ワ
I	イ	キ	シ	チ	ニ	ヒ	ミ	リ	ギ	ジ	ビ	ピ		
U	ウ	ク	ス	ツ	ヌ	フ	ム	ル	グ	ズ	ブ	プ	ユ	
E	エ	ケ	セ	テ	ネ	ヘ	メ	レ	ゲ	ゼ	ベ	ペ		
O	オ	コ	ソ	ト	ノ	ホ	モ	ロ	ゴ	ゾ	ボ	ポ	ヨ	ヲ
n					ン									

音節: カタカナ (71個)

音素: アルファベット (20個)

【図 1 6】



【図 17】

MIDI符号データベース(男性)

[A]	[I]	[U]	[E]	[O]	[K]	[S]	[T]	[N]	[H]
F5	A#6	D3	A#3	F#4	F4	B6	D#4	C4	D4
E5	C#3	C#3	A3	F4	A#3	A6	F3	D#3	A3
D#5	A#2	A#2	G#3	E4	C3	C4	D#3	G#2	F3
A4	D2	A2	E3	G#3	B2	B3	B2	G2	E3
G#4	C#2	D2	G#2	G3	A#2	F3	G#2	F#2	D#3
C2	C2	C#2	D2	F#3	G2	F2	F#2	F2	F#2
B1	B1	C2	C#2	C2	C#2	E2	E2	E2	E2
A#1	A#1	B1	A1	B1	A1	D#2	A1	A1	A1

(a)

(b)

【図 18】

MIDI符号データベース(女性)

[A]	[I]	[U]	[E]	[O]	[K]	[S]	[T]	[N]	[H]
B5	A6	E5	G6	B4	D#5	G#6	G#6	C3	F4
A#5	G#6	D#5	D#6	A#4	B4	E4	D6	B2	E4
F#5	G6	D#4	B4	A4	D#4	D#4	G#5	A#2	D#4
D#5	F#6	C#4	F4	G#4	C4	B3	F4	A2	B3
C#5	G3	D#3	E4	D4	A3	E3	E4	G#2	A#3
C5	F3	D3	D4	C#4	D#3	D3	D#4	F2	E3
B4	E3	C#3	C#4	A#3	D3	C#3	E3	C2	D#3
G#4	D#3	A#2	C4	A3	C#3	B2	G#2	A#1	B2

(a)

(b)

[M]	[R]	[G]	[Z]	[D]	[B]	[P]	[Y]	[W]	[n]
B3	B3	C#4	C#4	D4	C4	C4	C#4	E4	D3
D#3	D#3	B3	C4	F3	B3	A#3	B3	D4	B2
C3	C3	F3	C3	C3	F3	A3	A3	C#4	A2
G2	B2	C3	B2	G#2	D#3	F3	F3	B3	G#2
F#2	G#2	B2	G2	G2	C3	G2	C3	G#3	D2
F2	G2	G2	F#2	F2	B2	E2	G#2	G3	C#2
E2	F2	F#2	E2	E2	G2	D#2	D#2	C2	C2
D#2	A1	E2	D#2	A1	E2	A1	A1	A#1	B1

(c)

(d)

[M]	[R]	[G]	[Z]	[D]	[B]	[P]	[Y]	[W]	[n]
F#4	G#6	G6	A6	G#6	G4	E5	D#6	D#5	B5
C3	D#4	F#6	G#6	D#4	D#4	E4	F#4	D5	A#5
B2	C4	E4	G#5	C4	B3	D#4	E4	B4	C4
A#2	B3	D4	D3	B3	B2	D4	C4	A#4	A#3
G#2	C3	C4	C3	D3	A2	B3	B3	E4	D#3
F#2	B2	B3	A2	C3	G#2	A#3	E3	D#4	D3
F2	A2	D#3	G#2	G2	G2	E3	A2	D4	C3
C#2	G#2	B2	C2	C2	B1	D#3	G#2	C4	B2

(c)

(d)

【図 19】

各音素の五線譜化の例(男性)

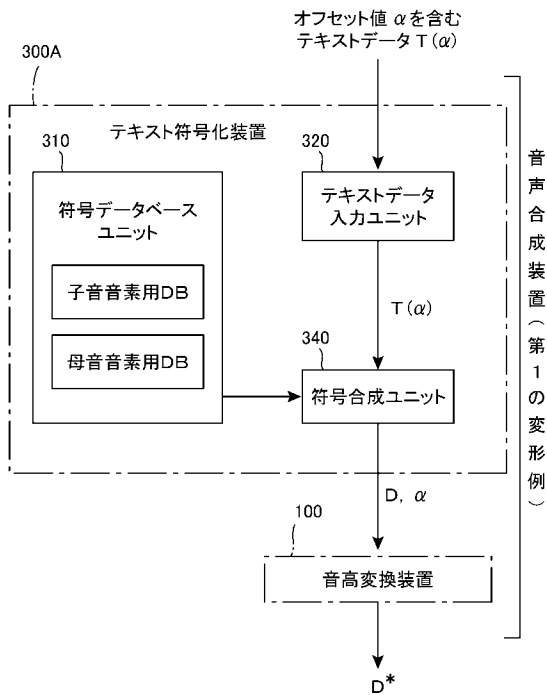


【図 20】

各音素の五線譜化の例(女性)



【図 2 1】



【図 2 2】

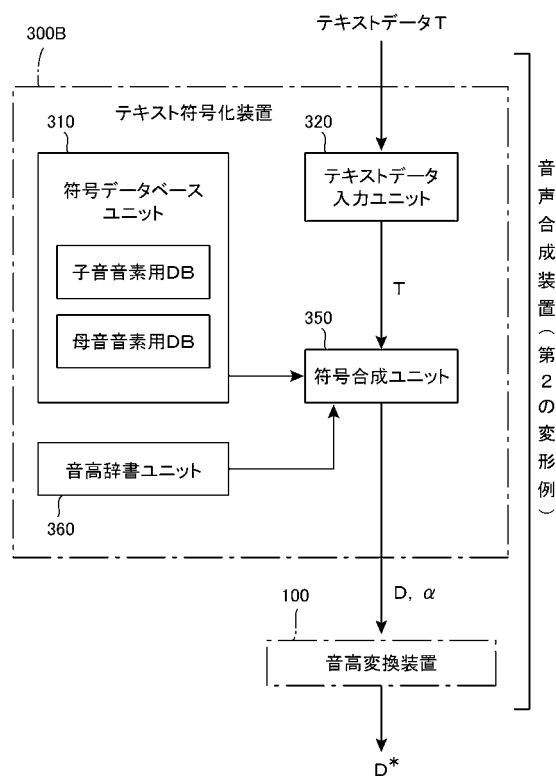
(a) 「コンニチワ」

(b) 「-3コ+3ン+6ニ+3チ-3ワ」

(c) $\underbrace{KO}_{-3} \underbrace{n}_{+3} \underbrace{NI}_{+6} \underbrace{TI}_{+3} \underbrace{WA}_{-3}$

- (d) 符号化音声データ $D(K), \alpha = -3$
 符号化音声データ $D(O), \alpha = -3$
 符号化音声データ $D(n), \alpha = +3$
 符号化音声データ $D(N), \alpha = +6$
 符号化音声データ $D(I), \alpha = +6$
 符号化音声データ $D(T), \alpha = +3$
 符号化音声データ $D(I), \alpha = +3$
 符号化音声データ $D(W), \alpha = -3$
 符号化音声データ $D(A), \alpha = -3$

【図 2 3】



【図 2 4】

360: 音高辞書ユニット

単語	オフセット値 α
オハヨウ	-1, +6, +4, -3
コンニチワ	-3, +3, +6, +3, -3
コンバンワ	-2, +2, +5, +2, -4
⋮	⋮
⋮	⋮

【図 2 5】

