

(12) 特許協力条約に基づいて公開された国際出願

(19) 世界知的所有権機関
国際事務局

(43) 国際公開日
2020年8月27日(27.08.2020)



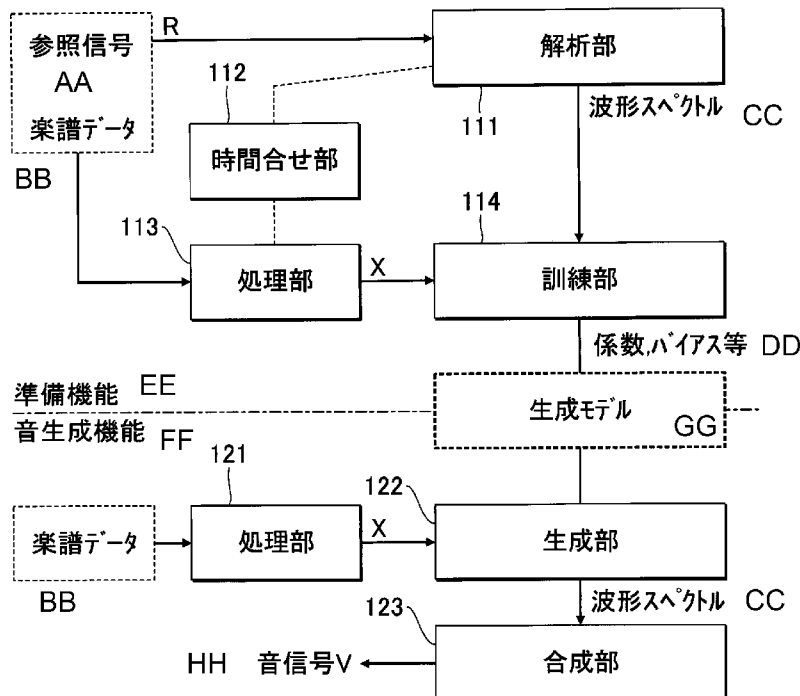
(10) 国際公開番号

WO 2020/171036 A1

- (51) 国際特許分類:
G10L 13/00 (2006.01) *G10H 7/08* (2006.01)
- (21) 国際出願番号: PCT/JP2020/006162
- (22) 国際出願日: 2020年2月18日(18.02.2020)
- (25) 国際出願の言語: 日本語
- (26) 国際公開の言語: 日本語
- (30) 優先権データ:
特願 2019-028684 2019年2月20日(20.02.2019) JP
- (71) 出願人: ヤマハ株式会社 (YAMAHA CORPORATION) [JP/JP]; 〒4308650 静岡県浜松市中区中沢町10番1号 Shizuoka (JP).
- (72) 発明者: 西村 方成 (NISHIMURA, Masanari); 〒4308650 静岡県浜松市中区中沢町10番1号 ヤマハ株式会社内 Shizuoka (JP).
- (74) 代理人: 大林 章, 外(OHBAYASHI, Akira et al.); 〒1130033 東京都文京区本郷2-15-13 お茶の水ウイングビル6階 Tokyo (JP).
- (81) 指定国(表示のない限り、全ての種類の国内保護が可能): AE, AG, AL, AM, AO, AT, AU, AZ, BA, BB, BG, BH, BN, BR, BW, BY, BZ, CA, CH, CL, CN, CO, CR, CU, CZ, DE, DJ, DK, DM, DO, DZ, EC, EE, EG, ES, FI, GB, GD, GE, GH, GM, GT, HN, HR, HU, ID, IL, IN, IR, IS, JO, JP, KE, KG, KH, KN, KP, KR, KW, KZ, LA, LC, LK, LR, LS, LU, LY, MA, MD, ME, MG, MK, MN, MW, MX, MY, MZ,

(54) Title: SOUND SIGNAL SYNTHESIS METHOD, GENERATIVE MODEL TRAINING METHOD, SOUND SIGNAL SYNTHESIS SYSTEM, AND PROGRAM

(54) 発明の名称: 音信号合成方法、生成モデルの訓練方法、音信号合成システムおよびプログラム



111 Analysis unit
112 Time setting unit
113, 121 Processing unit
114 Training unit
122 Generation unit
123 Synthesis unit
AA Reference signal
BB Musical score data
CC Waveform spectrum
DD Coefficient, bias, etc.
EE Preparation function
FF Sound generation function
GG Generative model
HH Sound signal

(57) Abstract: This sound signal synthesis method to be realized by a computer entails: setting a hot value for each piece of pitch name data that is specified by a pitch name of a sound signal to be synthesized, and that is part of pitch data which includes a plurality of pieces of pitch name data corresponding to a plurality of pitch names, the hot value corresponding to a shift from the pitch name of the pitch of the sound signal; using a generative model to estimate output data indicating the sound signal according to the set pitch data; and synthesizing the sound signal in accordance with the estimated output data.

NA, NG, NI, NO, NZ, OM, PA, PE, PG, PH, PL, PT,
QA, RO, RS, RU, RW, SA, SC, SD, SE, SG, SK, SL,
ST, SV, SY, TH, TJ, TM, TN, TR, TT, TZ, UA, UG,
US, UZ, VC, VN, WS, ZA, ZM, ZW.

- (84) 指定国(表示のない限り、全ての種類の広域保護が可能): ARIPO (BW, GH, GM, KE, LR, LS, MW, MZ, NA, RW, SD, SL, ST, SZ, TZ, UG, ZM, ZW), ユーラシア (AM, AZ, BY, KG, KZ, RU, TJ, TM), ヨーロッパ (AL, AT, BE, BG, CH, CY, CZ, DE, DK, EE, ES, FI, FR, GB, GR, HR, HU, IE, IS, IT, LT, LU, LV, MC, MK, MT, NL, NO, PL, PT, RO, RS, SE, SI, SK, SM, TR), OAPI (BF, BJ, CF, CG, CI, CM, GA, GN, GQ, GW, KM, ML, MR, NE, SN, TD, TG).

添付公開書類:

- 一 国際調査報告 (条約第21条(3))

(57) 要約: コンピュータにより実現される音信号合成方法は、複数の音名に対応する複数の音名データを含む音高データの、合成すべき音信号の音名で特定される音名データに、それぞれ、その音信号の音高のその音名からのずれに応じたホット値を設定し、生成モデルを用いて、設定された音高データに応じた音信号を示す出力データを推定し、推定された出力データに応じて、音信号を合成する。

明 細 書

発明の名称：

音信号合成方法、生成モデルの訓練方法、音信号合成システムおよびプログラム

技術分野

[0001] 本発明は、音信号を合成する音源技術に関する。

背景技術

[0002] 特許文献 1 に示すNSynth、または非特許文献 1 に示すNPSS (Neural Parametric Singing Synthesizer) など、ニューラルネットワーク（以下、「NN」と呼ぶ）を用いて、条件入力に応じた音波形を生成する音源（以下、「DNN (Deep Neural Network) 音源」と呼ぶ）が提案されている。NSynthは、エンベディング (embedding/埋込ベクトル) に応じて、サンプリング周期ごとに、音信号のサンプルを生成する。NPSSのTimbreモデルは、ピッチおよびタイミング情報に応じて、フレームごとに、音信号のスペクトルを生成する。

先行技術文献

特許文献

[0003] 特許文献1：米国特許第 1 0 0 6 8 5 5 7 号明細書

非特許文献

[0004] 非特許文献1：Merlijn Blaauw, Jordi Bonada, 「A Neural Parametric Singing Synthesizer Modeling Timbre and Expression from Natural Songs」, Appl. Sci. 2017, 7, 1313

発明の概要

発明が解決しようとする課題

[0005] 特許文献 1 のNsynthまたは非特許文献 1 のNPSSなどのDNN音源では、所望の 1 個の音階を指定する音高データにより、合成する音信号の音高を制御する。ピッチ包絡またはビブラートなど、音符などで指定される音階からの音高

の動的なずれを制御することについては、考慮されていない。

[0006] DNN音源の訓練フェーズでは、音高データを入力として、音信号または波形スペクトルを表す出力データを推定するNNを訓練する。DNN音源は、ビブラートのかかった音信号で訓練をすれば、ビブラートのかかった音信号を生成し、ピッチベンドがかかった音信号で訓練すれば、ピッチベンドのかかった音信号を生成する。しかしながら、ビブラートまたはピッチベンドのような、動的に変化する音高のずれ（ピッチベンド量）を、時間変化する数値で制御することはできない。

[0007] 本開示は、合成する音信号の動的なピッチ変化を、時間変化する数値で制御することを目的とする。

課題を解決するための手段

[0008] 本開示のひとつの態様に係る音信号合成方法は、合成すべき第1音信号の音高を表す第1音高データを生成し、第2音信号の音高を表す第2音高データと前記第2音信号との関係を学習した生成モデルを用いて、前記第1音高データに応じた前記第1音信号を示す出力データを推定するコンピュータにより実現される音信号合成方法であって、前記第1音高データは、相異なる音名に対応する複数の音名データを含み、前記第1音高データの生成においては、前記複数の音名データのうち、前記第1音信号の音高に対応する音名データを、当該音名データに対応する音名の基準音高と当該第1音信号の音高との相違に応じたホット値に設定する。

[0009] 本開示のひとつの態様に係る訓練方法は、合成すべき音信号の音高を表す音高データを用意し、前記音高データの入力に対して前記音信号を示す出力データが出力されるように生成モデルを訓練するコンピュータにより実現される生成モデルの訓練方法であって、前記音高データは、相異なる音名に対応する複数の音名データを含み、前記音高データを用意においては、前記複数の音名データのうち、前記音信号の音高に対応する音名データを、当該音名データに対応する音名の基準音高と当該音信号の音高との相違に応じたホット値に設定する。

[0010] 本開示のひとつの態様に係る音信号合成システムは、1以上のプロセッサと1以上のメモリとを具備する音信号合成システムであって、前記1以上のメモリは、第2音信号の音高を表す第2音高データと前記第2音信号との関係を学習した生成モデルを記憶し、前記1以上のプロセッサは、合成すべき第1音信号の音高を表す第1音高データを生成し、第2音信号の音高を表す第2音高データと前記第2音信号との関係を学習した生成モデルを用いて、前記第1音高データに応じた前記第1音信号を示す出力データを推定し、前記第1音高データは、相異なる音名に対応する複数の音名データを含み、前記1以上のプロセッサは、第1音高データの生成において、前記複数の音名データのうち、前記第1音信号の音高に対応する音名データを、当該音名データに対応する音名の基準音高と当該第1音信号の音高との相違に応じたホット値に設定する。

[0011] 本開示のひとつの態様に係るプログラムは、合成すべき第1音信号の音高を表す第1音高データを生成する処理部、および、第2音信号の音高を表す第2音高データと前記第2音信号との関係を学習した生成モデルを用いて、前記第1音高データに応じた前記第1音信号を示す出力データを推定する生成部としてコンピュータを機能させるプログラムであって、前記第1音高データは、相異なる音名に対応する複数の音名データを含み、前記第1音高データの生成においては、前記複数の音名データのうち、前記第1音信号の音高に対応する音名データを、当該音名データに対応する音名の基準音高と当該第1音信号の音高との相違に応じたホット値に設定する。

図面の簡単な説明

[0012] [図1]音信号合成システムのハードウェア構成を示すブロック図である。

[図2]音信号合成システムの機能構成を示すブロック図である。

[図3]音高データの説明図である。

[図4]訓練部と生成部による処理の説明図である。

[図5]1ホットレベル記法に従う音高データの説明図である。

[図6]準備処理のフローチャートである。

[図7]音生成処理のフローチャートである。

[図8]2 ホットレベル記法に従う音高データの説明図である。

[図9]4 ホットレベル記法に従う音高データの説明図である。

[図10]各音名と音信号の音高の近接度の変形例を示す図である。

発明を実施するための形態

[0013] A : 第1 実施形態

図1は、本開示の音信号合成システム100の構成を例示するブロック図である。音信号合成システム100は、制御装置11と記憶装置12と表示装置13と入力装置14と放音装置15とを具備するコンピュータシステムで実現される。音信号合成システム100は、例えば携帯電話機、スマートフォンまたはパーソナルコンピュータ等の情報端末である。音信号合成システム100は、単体の装置で実現されるほか、相互に別体で構成された複数の装置（例えばサーバクライアントシステム）でも実現される。

[0014] 制御装置11は、音信号合成システム100を構成する各要素を制御する単数または複数のプロセッサである。具体的には、例えばCPU（Central Processing Unit）、SPU（Sound Processing Unit）、DSP（Digital Signal Processor）、FPGA（Field Programmable Gate Array）、またはASIC（Application Specific Integrated Circuit）等の1種類以上のプロセッサにより、制御装置11が構成される。制御装置11は、合成音の波形を表す時間領域の音信号Vを生成する。

[0015] 記憶装置12は、制御装置11が実行するプログラムと制御装置11が使用する各種のデータとを記憶する単数または複数のメモリである。記憶装置12は、例えば磁気記録媒体もしくは半導体記録媒体等の公知の記録媒体、または、複数種の記録媒体の組合せで構成される。なお、音信号合成システム100とは別体の記憶装置12（例えばクラウドストレージ）を用意し、移動体通信網またはインターネット等の通信網を介して制御装置11が記憶装置12に対する書込および読出を実行してもよい。すなわち、記憶装置12は音信号合成システム100から省略されてもよい。

- [0016] 表示装置 1 3 は、制御装置 1 1 が実行したプログラムの演算結果を表示する。表示装置 1 3 は、例えばディスプレイである。表示装置 1 3 は音信号合成システム 1 0 0 から省略されてもよい。
- [0017] 入力装置 1 4 は、ユーザの入力を受け付ける。入力装置 1 4 は、例えばタッチパネルである。入力装置 1 4 は音信号合成システム 1 0 0 から省略されてもよい。
- [0018] 放音装置 1 5 は、制御装置 1 1 が生成した音信号 V が表す音声を再生する。放音装置 1 5 は、例えばスピーカまたはヘッドホンである。なお、制御装置 1 1 が生成した音信号 V をデジタルからアナログに変換する D/A 変換器と音信号 V を増幅する増幅器とについては図示を便宜的に省略した。また、図 1 では、放音装置 1 5 を音信号合成システム 1 0 0 に搭載した構成を例示したが、音信号合成システム 1 0 0 とは別体の放音装置 1 5 を音信号合成システム 1 0 0 に有線または無線で接続してもよい。
- [0019] 図 2 は、音信号合成システム 1 0 0 の機能構成を示すブロック図である。制御装置 1 1 は、記憶装置 1 2 に記憶されたプログラムを実行することで、生成モデルを用いて、歌手の歌唱音または楽器の演奏音などの音波形を表す時間領域の音信号 V を生成する音生成機能（処理部 1 2 1、生成部 1 2 2、および合成部 1 2 3）を実現する。また、制御装置 1 1 は、記憶装置 1 2 に記憶されたプログラムを実行することで、音信号 V の生成に用いられる生成モデルの準備を行う準備機能（解析部 1 1 1、時間合せ部 1 1 2、処理部 1 1 3、および訓練部 1 1 4）を実現する。なお、複数の装置の集合（すなわちシステム）で制御装置 1 1 の機能を実現してもよいし、制御装置 1 1 の機能の一部または全部を専用の電子回路（例えば信号処理回路）で実現してもよい。
- [0020] まず、音高データ X1 と、音高データ X1 に応じて出力データを生成する生成モデルと、生成モデルの訓練に用いる参照信号 R について説明する。
- [0021] 音高データ X1 は、音信号（参照信号 R または音信号 V）の音高（以下「目標音高」という）P を示すデータである。図 3 は、音高データ X1 の例である

。音高データX1は、相異なる音名（…”G # 3”, ”A 3”, ”A # 3”, ”B 3”, ”C 4”, ”C # 4”, ”D 4”, …）に対応する複数（M個）の音名データを有する（Mは2以上の自然数）。なお、音名を表す記号（C, D, E, …）が共通する音名であってもオクターブが相違する音名は別個の音名として区別する。

[0022] 音高データX1を構成するM個の音名データのうち、目標音高Pに対応する1個の音名データ（以下「有効音名データ」という）は、当該有効音名データの音名に対応する所定の音高（以下「基準音高」という）Qと目標音高Pとの音高差（ずれ）に応じたずれ値に設定される。ずれ値は、ホット値（Hot Value）の一例である。1個の音名に対応する基準音高Qは、当該音名に対応する標準的な音高（ピッチ）である。他方、音高データX1を構成するM個の音名データのうち、有効音名データ以外の（M-1）個の音名データは、目標音高Pに関与しないことを示す定数値（例えば0）に設定される。目標音高Pに関与しないことを示す定数値は、コールド値の一例である。以上の説明から理解される通り、音高データX1は、音信号（参照信号Rまたは音信号V）の目標音高Pに対応する音名と、当該音名の基準音高Qに対する目標音高Pのずれ値との双方を指定する。

[0023] 生成モデルは、音高データX1を含む制御データXに応じて、音信号Vの波形スペクトル（例えば、メルスペクトログラム、または基本周波数などの特徴量）の時系列を生成する統計的モデルである。制御データXは、合成されるべき音信号の条件を指定するデータである。生成モデルの生成特性は、記憶装置12に記憶された複数の変数（係数およびバイアスなど）により規定される。統計的モデルは、波形スペクトルを推定するニューラルネットワークである。そのニューラルネットワークは、例えば、WaveNet(TM)のような、音信号Vの過去の複数のサンプルに基づいて、現在のサンプルの確率密度分布を推定する回帰的なタイプでもよい。また、そのアルゴリズムも任意であり、例えば、CNN（Convolutional Neural Network）タイプでもRNN（Recurrent Neural Network）タイプでよいし、その組み合わせでもよい。さらに、LS

TM (Long Short-Term Memory) またはATTENTIONなどの付加的要素を備えるタイプでもよい。生成モデルの複数の変数は、後述する準備機能による訓練データを用いた訓練により確立される。複数の変数が確立された生成モデルは、後述する音生成機能において音信号Vの生成に使用される。

[0024] 記憶装置12は、生成モデルの訓練のために、ある楽譜をプレイヤーが演奏した時間領域の波形を示す音信号（以下、「参照信号」と呼ぶ）Rと、当該楽譜を表す楽譜データとが複数記録されている。各楽譜データは音符の時系列を含む。各楽譜データに対応する参照信号Rは、当該楽譜データが表す楽譜の音符の系列に対応する部分波形の時系列を含む。各参照信号Rは、サンプリング周期（例えば、48kHz）ごとのサンプルの時系列で構成され、音の波形を表す時間領域の信号である。演奏は、人間による楽器の演奏に限らず、歌手による歌唱、または楽器の自動演奏であってもよい。機械学習で良い音を生成するためには、一般的に十分な個数の訓練データが要求されるので、ターゲットとする楽器またはプレイヤーなどについて、多数の演奏の音信号を事前に収録し、参照信号Rとして記憶装置12に記憶しておく。

[0025] 図2の上段に図示される準備機能について説明する。解析部111は、複数の楽譜にそれぞれ対応する複数の参照信号Rの各々について、時間軸上のフレームごとに周波数領域のスペクトル（以下「波形スペクトル」と呼ぶ）を算定する。参照信号Rの波形スペクトルの算定には、例えば離散フーリエ変換等の公知の周波数解析が用いられる。波形スペクトルには、基本周波数などの音響的な特徴量も含まれる。

[0026] 時間合せ部112は、解析部111で得られた波形スペクトル等の情報に基づき、各参照信号Rに対応する楽譜データにおける複数の発音単位の開始時点と終了時点とを、参照信号Rにおけるその発音単位に対応する部分波形の開始時点と終了時点とに揃える。ここで、発音単位は、例えば、音高と発音期間とが指定された1個の音符である。なお、1個の音符を、音色等の波形の特徴が変化するポイントで分割して、複数の発音単位に分けてもよい。

[0027] 処理部113は、各参照信号Rに時間が揃えられた楽譜データの各発音単

位の情報に基づき、フレームを単位とする時刻 t ごとに、参照信号 R のうち当該時刻 t の部分波形に対応する制御データ X を生成する。処理部 113 が生成する制御データ X は、前述の通り、参照信号 R の条件を指定する。

[0028] 制御データ X は、図 4 に例示される通り、音高データ $X1$ と開始停止データ $X2$ とコンテキストデータ $X3$ とを含む。音高データ $X1$ は、参照信号 R の部分波形における目標音高 P を表す。開始停止データ $X2$ は、各部分波形の開始期間（アタック）と終了期間（リリース）とを表す。1 個の音符に相当する部分波形内の 1 個のフレームのコンテキストデータ $X3$ は、当該音符と前後の音符との音高差、または楽譜内における当該音符の相対的な位置を表す情報など、複数の発音単位との関係（すなわちコンテキスト）を表す。制御データ X には、さらに、楽器、歌手または奏法など、その他の情報を含んでいてもよい。

[0029] 前述の通り、音高データ $X1$ を構成する M 個の音名データのうち、音信号（参照信号 R または音信号 V ）の目標音高 P に対応する 1 個の有効音名データは、当該音名に対応する基準音高 Q に対する目標音高 P の音高差に応じたずれ値に設定される。この記法に従う音高データ $X1$ を、1 ホットレベル記法の音高データ $X1$ と呼ぶ。処理部 113（制御装置 11）は、音高データ $X1$ の M 個の音名データのうち、参照信号 R の目標音高 P に対応する 1 個の有効音名データを、当該音名に対応する基準音高 Q と目標音高 P との音高差に応じたずれ値に設定する。図 5 に、その設定例を示す。

[0030] 図 5 の上段には、時間軸（横軸）と音高軸（縦軸）とが設定された 2 次元平面に、楽譜データが表す楽譜を構成する音符の時系列と、当該楽譜を演奏した演奏音の音高（目標音高 P ）とが図示されている。図 5 の例では、音符 $F \#$ 、音符 F 、休符、音符 F 、音符 $F \#$ 、音符 F が順に演奏される。図 5 における目標音高 P は、例えば音高が連続的に変化する楽器から発音される演奏音の音高である。

[0031] 図 5 に例示される通り、音高軸は、相異なる音名に対応する複数の範囲（以下「単位範囲」という） U に区画される。各音名に対応する基準音高 Q は

、例えば当該音名に対応する単位範囲Uの中点に相当する。例えば、音名F #に対応する基準音高Q (F#)は、当該音名F #に対応する単位範囲U (F#)の中点である。図5から理解される通り、目標音高Pが各音符の基準音高Qに近付くように楽曲が演奏される。各音名に対応する基準音高Qは音高軸上に離散的に設定されるのに対し、目標音高Pは時間の経過とともに連続的に変化する。したがって、目標音高Pは基準音高Qに対してずれている。

[0032] 図5の中段は、音高データX1の各音名データが表す数値の時間的な変化を表すグラフである。図5の中段における縦軸の数値0は、各音名に対応する基準音高Qを意味する。目標音高Pが1個の音名の単位範囲U内にある場合、音高データX1のM個の音名データのうち当該音名に対応する1個の音名データが有効音名データとして選択され、当該有効音名データが基準音高Qに対するずれ値に設定される。

[0033] 有効音名データが表すずれ値は、当該有効音名データの音名に対応する基準音高Q (= 0) に対する目標音高Pの相対値である。1個の音名に対応する単位範囲Uの幅は100セント（半音分）であるから、基準音高Qに対する目標音高Pの音高差は±50セントの範囲に収まる。一方、有効音名データに設定されるずれ値は、0から1の範囲の何れかの値をとる。音高差とずれ値との対応付けは任意であるが、例えば、ずれ値の0から1の範囲が、音高差の-50セントから+50セントに対応付けられる。例えば、ずれ値の0が音高差の-50セントに対応し、ずれ値の0.5が音高差の0セントに対応し、ずれ値の1が音高差の+50セントに対応する。

[0034] 図5に例示される通り、時間軸上の時刻t1において、目標音高Pは、音名F #に対応する単位範囲U (F#)内にあり、かつ、当該音名F #に対応する基準音高Q (F#)との音高差は+40セントである。したがって、時刻t1においては、音高データX1のM個の音名データのうち、音名F #に対応する1個の有効音名データが、+40セントの音高差に対応するずれ値0.9に設定され、残余の(M-1)個の音名データは0（コールド値）に設定される。

[0035] また、時刻t2において、目標音高Pは、音名Fに対応する単位範囲U (F)

内にあり、かつ、当該音名 F に対応する基準音高 Q (F) との音高差は + 20 セントである。したがって、時刻 t_2 においては、音高データ X1 の M 個の音名データのうち、音名 F に対応する 1 個の有効音名データが、+ 20 セントの音高差に対応するずれ値 0.7 に設定される。

[0036] なお、音高差とずれ値との対応は上記に限らず、例えば、ずれ値の 0.2 から 1 の範囲を、音高差の - 50 セントから + 50 セントの範囲に対応付けてもよい。例えば、ずれ値の 0.2 が音高差の - 50 セントに対応し、ずれ値の 0.6 が音高差の 0 セントに対応し、ずれ値の 1 が音高差の + 50 セントに対応する。或いは、ずれ値の正負と音高差の正負との関係を反転し、ずれ値の 0.2 から 1 の範囲を、音高差の + 50 セントから - 50 セントに対応付けてもよい。

[0037] 解析部 111 および処理部 113 による処理の結果、複数の参照信号 R と複数の楽譜データとから、生成モデルを訓練するための複数の発音単位データが準備される。各発音単位データは、制御データ X と波形スペクトルとのセットである。複数の発音単位データは、訓練部 114 による訓練に先立ち、生成モデルの訓練のための訓練データと、生成モデルのテストのためのテストデータとに分けられる。複数の発音単位データの大部分を訓練データとし、一部をテストデータにする。訓練データによる訓練は、複数の発音単位データをフレームの所定個ごとにバッチとして分割し、バッチ単位で全バッチにわたり順番に行われる。

[0038] 訓練部 114 は、図 4 の上段に例示するように、訓練データを受け取り、その各バッチの複数の発音単位の波形スペクトルと制御データ X とを順番に用いて生成モデルを訓練する。生成モデルは、フレーム（時刻 t ）ごとに、波形スペクトルを表す出力データを推定する。出力データは、波形スペクトルを構成する複数の成分の各々の確率密度分布を示すデータであっても良いし、各成分の値であってもよい。訓練部 114 は、1 バッチ分の各発音単位データにおける制御データ X を生成モデルに入力することで、その制御データ X に対応する出力データの時系列を推定する。訓練部 114 は、推定され

た出力データと訓練データのうち対応する波形スペクトル（すなわち正解値）とに基づいて損失関数 L （１バッチ分の累算値）を計算する。そして、訓練部１１４は、その損失関数 L が最小化されるように生成モデルの複数の変数を最適化する。例えば、損失関数 L としては、出力データが確率密度分布である場合にはクロスエントロピー関数などが使用され、出力データが波形スペクトルの値である場合には二乗誤差関数などが使用される。訓練部１１４は、訓練データを利用した以上の訓練を、テストデータについて算出される損失関数 L の値が十分に小さくなるか、或いは、相前後する損失関数 L の変化が十分に小さくなるまで繰り返し行う。こうして確立された生成モデルは、各時刻 t における制御データ X と、参照信号 R のうち当該時刻 t に対応する波形スペクトルとの間に潜在する関係を習得している。この生成モデルを用いることで、生成部１２２は、未知の音信号 V の制御データ X についても、品質の良い波形スペクトルを生成できる。

[0039] 図６は、準備処理のフローチャートである。準備処理は、例えば音信号合成システム１００の利用者からの指示を契機として開始される。

[0040] 準備処理が開始されると、制御装置１１（解析部１１１）は、複数の参照信号 R の各々から各部分波形の波形スペクトルを生成する（Sa1）。次に、制御装置１１（時間合せ部１１２および処理部１１３）は、その部分波形に対応する楽譜データから、その部分波形に対応する発音単位の音高データ $X1$ を含む制御データ X を作成する（Sa2）。制御装置１１（訓練部１１４）は、各発音単位の各時刻 t の制御データ X と当該発音単位に対応する波形スペクトルとを用いて生成モデルを訓練し、生成モデルの複数の変数を確立する（Sa3）。

[0041] ここでは、各音名の基準音高 Q に対するずれ値を示す音高データ $X1$ を含む制御データ X を入力として生成モデルが訓練される。したがって、その訓練により確立された生成モデルは、制御データ X が示す音高のずれ値と、音信号（参照信号 R ）の波形スペクトルとの間に潜在する関係を習得する。これにより、音名とずれ値とを指定する音高データ $X1$ を含む制御データを生成モ

デルに入力すれば、その指定された音名とずれに応じた音高の音信号Vを生成できる。

[0042] 本願の発明者は、対比例として、音信号の音名を示す従来のone-hotの音高データと、当該音名の基準音高Qに対する音信号の音高ずれを示すベンドデータとをパラレルに含む制御データを入力として生成モデルを訓練した。しかしながら、その訓練で確立された生成モデルを用いた場合、生成される音信号の音高は、音高データが示す音高に追従するものの、ベンドデータが示すずれに安定して追従しなかった。これは、生成モデルで生成される音信号の1個の特徴量である音高を、音高データおよびベンドデータという別個の2種類のデータで制御しようとしたため、と考察される。

[0043] 続いて、図2の下段に図示される、生成モデルを用いて音信号Vを生成する音生成機能について説明する。処理部121は、処理部113と同様に、再生すべき楽譜データが表す一連の発音単位に基づき、制御データXを生成して生成部122に出力する。制御データXは、楽譜データの各時刻tにおける発音単位の条件（すなわち、合成されるべき音信号Vの条件）を表す。具体的には、制御データXは、音高データX1と開始停止データX2とコンテキストデータX3とを含む。処理部113が生成する音高データX1が、参照信号Rの目標音高Pを表すのに対し、処理部121が生成する音高データX1は、合成すべき音信号Vの目標音高Pを表す。ただし、処理部113が実行する処理と処理部121が実行する処理とは実質的に共通し、処理部113が生成する音高データX1の形式と処理部121が生成する音高データX1の形式とは共通する。なお、制御データXは、さらに、楽器、歌手または奏法など、その他の情報を含んでもよい。

[0044] 生成部122は、図4の下段に例示するように、複数の変数が確立された生成モデルを用いて、制御データXに応じた波形スペクトルの時系列を生成する。生成部122は、生成モデルを用いて、フレームごと（時刻t）に、制御データXに応じた波形スペクトルを示す出力データを推定する。推定される出力データが、波形スペクトルを構成する複数の成分の各々の確率密度分

布を表す場合、生成部 1 2 2 は、その成分の確率密度分布に従う乱数を生成して、当該乱数を波形スペクトルの成分値として出力する。推定される出力データが複数の成分の値を表す場合は、その成分値を出力する。

[0045] 合成部 1 2 3 は、周波数領域の波形スペクトルの時系列を受け取り、その波形スペクトルの時系列に応じた時間領域の音信号 V を合成する。合成部 1 2 3 は、いわゆるボコーダである。例えば、合成部 1 2 3 は、波形スペクトルから最小位相スペクトルを求めて、それら波形スペクトルと位相スペクトルとに対して逆フーリエ変換を実行することで音信号 V を合成する。或いは、波形スペクトルと音信号 V の間に潜在する関係を学習したニューラルボコーダを用いて、波形スペクトルから音信号 V を直接的に合成する。

[0046] 図 7 は、各発音単位の音生成処理のフローチャートである。この音生成処理は、例えば音信号合成システム 1 0 0 の利用者からの指示を契機として、時刻 t 毎に、当該時刻に対応するフレームの音信号 V を生成するために実行される。時刻 t の進行速度は、実時間の進行速度と同じ位でも良いし、速いあるいは遅い、つまり、実時間と異なってもよい。

[0047] ある時刻 t の音生成処理が開始されると、制御装置 1 1（処理部 1 2 1）は、楽譜データに基づいて、その時刻 t の制御データ X を生成する（Sb1）。次に、制御装置 1 1（生成部 1 2 2）は、生成モデルを用いて、生成された制御データ X に応じた、その時刻 t の音信号 V の波形スペクトルを生成する（Sb2）。次に、制御装置 1 1（合成部 1 2 3）は、生成された波形スペクトルに応じて、その時刻 t に対応するフレームの音信号 V を合成する（Sb3）。以上の処理が、楽譜データの各時刻 t について順次行われることで、楽譜データに対応する音信号 V が生成される。

[0048] 第 1 実施形態においては、1 個の音高データ X1 が、合成されるべき音信号 V の目標音高 P と当該音名の基準音高 Q との音高差に応じたずれ値とを指定する。そして、生成部 1 2 2 は、この音高データ X1 を含む制御データ X を入力とする生成モデルを用いて、音高データ X1 で指定された音名とずれ値とに応じた音高の音信号 V を生成する。したがって、生成される音信号 V の音高

は、音高データ X1 が指定する音名と、その音名の基準音高 Q からのずれ値との変化に良く追従する。例えば、音高データ X1 が示すずれ値を動的に変化させることにより、生成する音信号 V に、例えばビブラートまたはピッチベンドなどの、動的なピッチ変化を付与できる。

[0049] B：第 2 実施形態

第 2 実施形態では、生成モデルへの入力として、第 1 実施形態の 1 ホットレベル記法の音高データ X1 の代わりに、図 8 に例示する 2 ホットレベル記法の音高データ X1 を用いる。第 2 実施形態の音信号合成システム 100 の構成、及び、制御装置 11 の機能構成は、基本的に第 1 実施形態と同じである。

[0050] 2 ホットレベル記法の音高データ X1 においては、相異なる音名に対応する M 個の音名データのうち、音信号（参照信号 R または音信号 V）の目標音高 P に対応する 2 個の有効音名データの各々に、当該有効音名データの音名に対応する基準音高 Q と目標音高 P との音高差に応じたホット値が設定される。処理部 113 または処理部 121（制御装置 11）は、音高データ X1 の M 個の音名データのうち、音信号（参照信号 R または音信号 V）の目標音高 P を挟む 2 個の基準音高 Q の各々に対応する音名データを有効音名データとして選択し、2 個の有効音名データの各々を、目標音高 P と当該音名データの音名に対応する基準音高 Q との近接度（ホット値の一例）に設定する。すなわち、2 ホットレベル記法とは、音高データ X1 を構成する M 個の音名データのうちの 2 個の有効音名データをホット値（近接度）に設定し、残余の (M - 2) 個の音名データをコールド値（例えば 0）に設定する記述方法である。

[0051] 図 8 の上段には、時間軸（横軸）と音高軸（縦軸）とが設定された 2 次元平面に、楽譜データが表す楽譜を構成する音符の時系列と、当該楽譜を演奏した演奏音の音高（目標音高 P）とが図示されている。図 8 から分かるように、音高データ X1 を構成する M 個の音名データのうち、目標音高 P に 1 番近い基準音高 Q に対応する音名の音名データと、2 番目に近い基準音高 Q に対応する音名の音名データとが、有効音名データとして選択される。

[0052] 図 8 の中段に、各音名データの音名に対応する基準音高 Q と目標音高 P と

の近接度を示す。ここで、近接度は、0から1の範囲の何れかの値を取る。具体的には、目標音高Pとある音名の基準音高Qとが一致するときに近接度は1である。また、目標音高Pとその音名の基準音高Qとの音高差がxセントのとき、近接度は $(100 - x) / 100$ である。すなわち、目標音高Pと基準音高Qとの音高差が大きいほど近接度は小さい数値となる。例えば、目標音高Pがある音名の基準音高Qから半音以上離れると、近接度は0となる。

[0053] 図8の時刻t3においては、目標音高Pが、音名Gに対応する基準音高Q(G)と音名F#に対応する基準音高Q(F#)との間に位置する。したがって、音高データX1のM個の音名データのうち、音名Gに対応する音名データと音名F#に対応する音名データとが有効音名データとして選択される。時刻t3において、基準音高Q(G)と目標音高Pとの音高差は50セントであるから、音名Gに対応する有効音名データは、近接度0.5に設定される。また、時刻t3において、基準音高Q(F#)と目標音高Pとの音高差も50セントであるから、音名F#に対応する有効音名データも、近接度0.5に設定される。以上の通り、第2実施形態の処理部113または処理部121は、時刻t3において、音高データX1を構成するM個の音名データのうち、音名Gに対応する有効音名データを0.5に設定し、音名F#に対応する有効音名データを0.5に設定し、残余の(M-2)個の音名データを0(コールド値)に設定する。

[0054] 他方、図8の時刻t4においては、目標音高Pが、音名Fに対応する基準音高Q(F)と音名F#に対応する基準音高Q(F#)との間に位置する。したがって、音高データX1のM個の音名データのうち音名Fに対応する音名データと音名F#に対応する音名データとが有効音名データとして選択される。時刻t4において、基準音高Q(F)と目標音高Pとの音高差は80セントであるから、音名Fに対応する有効音名データは、近接度0.2に設定される。また、時刻t4において、基準音高Q(F#)と目標音高Pとの音高差は20セントであるから、音名F#に対応する有効音名データは、近接度0.8に設定される。

以上の通り、第2実施形態の処理部113または処理部121は、時刻 t_4 において、音高データ X_1 を構成する M 個の音名データのうち、音名 F に対応する有効音名データを0.2に設定し、音名 $F\#$ に対応する有効音名データを0.8に設定し、残余の $(M-2)$ 個の音名データを0（コールド値）に設定する。

[0055] 訓練部114は、2ホットレベル表記の音高データ X_1 を含む制御データ X を入力として、その制御データ X に対応する波形スペクトルを示す出力データを生成するよう、生成モデルを訓練する。複数の変数が確立された生成モデルは、複数の発音単位データにおける各制御データ X と、参照信号 R の波形スペクトルとの間に潜在する関係を習得している。

[0056] 生成部122は、確立された生成モデルを用いて、2ホットレベル表記の音高データ X_1 を含む制御データ X に応じた波形スペクトルを時刻 t 毎に生成する。合成部123は、第1実施形態と同様に、生成部122が生成した波形スペクトルの時系列に応じた時間領域の音信号 V を合成する。

[0057] 第2実施形態においては、生成モデルを用いて、2ホットレベル表記の音高データ X_1 が表す目標音高 P の変化によく追従する音信号 V を生成できる。

[0058] C：第3実施形態

第2実施形態においては、目標音高 P に対応する2個の有効音名データがホット値に設定されるが、音高データ X_1 を構成する M 個の音名データのうちホット値に設定される有効音名データの個数は任意である。第3実施形態においては、第2実施形態の2ホットレベル記法の音高データ X_1 の代わりに、図9に例示する4ホットレベル記法の音高データ X_1 を生成モデルに対する入力として利用する。第3実施形態の音信号合成システム100の構成、及び、制御装置11の機能構成は、基本的に第1実施形態および第2実施形態と同じである。

[0059] 4ホットレベル記法の音高データ X_1 においては、相異なる音名に対応する M 個の音名データのうち、音信号（参照信号 R または音信号 V ）の目標音高 P に対応する4個の音名データが有効音名データとして選択される。具体的

には、目標音高 P を挟む 2 個の基準音高 Q の各々に対応する 2 個の音名データと、当該 2 個の音名データに隣合う 2 個の音名データとが、有効音名データとして選択される。目標音高 P に近い 4 個の音名データが有効音名データとして選択されると換言してもよい。4 個の有効音名データの各々は、目標音高 P と当該有効音名データの音名に対応する基準音高 Q との近接度（ホット値）に設定される。すなわち、4 ホットレベル記法とは、音高データ X1 を構成する M 個の音名データのうちの 4 個の音名データをホット値（近接度）に設定し、残余の (M - 4) 個の音名データをコールド値（例えば 0）に設定する記述方法である。処理部 113 または処理部 121（制御装置 11）は、以上に説明した音高データ X1 を生成する。

[0060] 図 9 の上段には、第 2 実施形態と同様に、時間軸（横軸）と音高軸（縦軸）とが設定された 2 次元平面に、楽譜データが表す楽譜を構成する音符の時系列と、当該楽譜を演奏した演奏音の音高（目標音高 P）とが図示されている。図 9 から理解される通り、音高データ X1 を構成する M 個の音名データのうち、目標音高 P に近い 4 個の基準音高 Q にそれぞれ対応する 4 個の音名データが有効音名データとして選択される。

[0061] 図 9 の中段に、各音名データの音名に対応する基準音高 Q と目標音高 P との近接度を示す。ここで、近接度は、第 2 実施形態と同様に、0 から 1 の範囲の何れかの値を取る。具体的には、目標音高 P とある音名の基準音高 Q とが一致するときに近接度は 1 である。また、目標音高 P とその音名の基準音高 Q との音高差が x セントのとき、近接度は $(200 - x) / 200$ である。すなわち、第 2 実施形態と同様に、目標音高 P と基準音高 Q との音高差が大きいほど近接度は小さい数値となる。例えば、目標音高 P がある音名の基準音高 Q から全音以上離れると、近接度は 0 となる。

[0062] 図 9 の時刻 t5 においては、目標音高 P が、音名 G に対応する基準音高 Q (G) と音名 F # に対応する基準音高 Q (F#) との間に位置する。したがって、音高データ X1 の M 個の音名データのうち、音名 G および音名 F # と、音名 G の高位側に隣合う音名 G # と、音名 F # の低位側に隣合う音名 F とに対応する 4

個の音名データが、有効音名データとして選択される。時刻 t_5 において、基準音高 $Q(G)$ と目標音高 P との音高差は 50 セントであるから、音名 G に対応する有効音名データは近接度 0.75 に設定される。同様に、基準音高 $Q(F\#)$ と目標音高 P との音高差は 50 セントであるから、音名 $F\#$ に対応する有効音名データは近接度 0.75 に設定される。また、基準音高 $Q(F)$ と目標音高 P との音高差は 150 セントであるから、音名 F に対応する有効音名データは近接度 0.25 に設定される。同様に、基準音高 $Q(G\#)$ と目標音高 P との音高差は 150 セントであるから、音名 $G\#$ に対応する有効音名データは近接度 0.25 に設定される。以上の通り、第 3 実施形態の処理部 113 または処理部 121 は、時刻 t_5 において、音高データ X_1 を構成する M 個の音名データのうち、音名 G および音名 $F\#$ に対応する 2 個の有効音名データを 0.5 に設定し、音名 F および音名 $G\#$ に対応する 2 個の有効音名データを 0.25 に設定し、残余の $(M-4)$ 個の音名データを 0 (コールド値) に設定する。

[0063] 時刻 t_6 においては、目標音高 P が、音名 $F\#$ に対応する基準音高 $Q(F\#)$ と音名 F に対応する基準音高 $Q(F)$ と音間に位置する。したがって、音高データ X_1 の M 個の音名データのうち、音名 $F\#$ および音名 F と、音名 $F\#$ の高位側に隣合う音名 G と、音名 F の低位側に隣合う音名 E とに対応する 4 個の音名データが、有効音名データとして選択される。時刻 t_6 において、基準音高 $Q(F\#)$ と目標音高 P との音高差は 25 セントであるから、音名 $F\#$ に対応する有効音名データは近接度 0.875 に設定される。基準音高 $Q(F)$ と目標音高 P との音高差は 75 セントであるから、音名 F に対応する有効音名データは近接度 0.625 に設定される。基準音高 $Q(G)$ と目標音高 P との音高差は 125 セントであるから、音名 G に対応する有効音名データは近接度 0.375 に設定される。また、基準音高 $Q(E)$ と目標音高 P との音高差は 175 セントであるから、音名 E に対応する有効音名データは近接度 0.125 に設定される。以上の通り、第 3 実施形態の処理部 113 または処理部 121 は、時刻 t_6 において、音高データ X_1 を構成する M 個の音名データのうち、音名

F #に対応する有効音名データを0. 8 7 5に設定し、音名Fに対応する有効音名データを0. 6 2 5に設定し、音名Gに対応する有効音名データを0. 3 7 5に設定し、音名Eに対応する有効音名データを0. 1 2 5に設定し、残余の(M-4)個の音名データを0 (コールド値) に設定する。

[0064] 訓練部1 1 4は、4 ホットレベル表記の音高データX1を含む制御データXを入力として、その制御データXに対応する波形スペクトルを示す出力データを生成するよう、生成モデルを訓練する。複数の変数が確立された生成モデルは、複数の発音単位データにおける各制御データXと、参照信号Rの波形スペクトルとの間に潜在する関係を習得している。

[0065] 生成部1 2 2は、確立された生成モデルを用いて、2 ホットレベル表記の音高データX1を含む制御データXに応じた波形スペクトルを時刻t毎に生成する。合成部1 2 3は、第1実施形態と同様に、生成部1 2 2が生成した波形スペクトルの時系列に応じた時間領域の音信号Vを合成する。

[0066] 第3実施形態においては、生成モデルを用いて、4 ホットレベル表記の音高データX1が表す目標音高Pの変化によく追従する音信号Vを生成できる。

[0067] 第1実施形態に例示した1 ホットレベル記法と、第2実施形態に例示した2 ホットレベル記法と、第3実施形態に例示した4 ホットレベル記法とは、音高データX1内の有効音名データの個数をNとしたNホットレベル記法として一般化される(Nは1以上の自然数)。Nホットレベル記法においては、音高データX1を構成するM個の音名データのうち、目標音高Pに対応するN個の有効音名データが、当該音名の基準音高Qと目標音高Pとの音高差に応じたホット値(ずれ値または近接度)に設定され、残余の(M-N)個の音名データがコールド値(例えば0)に設定される。目標音高Pとある音名の基準音高Qとの音高差がxセントであるとき、近接度は $(50 \times N - x) / 50 \times N$ と表現される。ただし、近接度を算定するための演算式は以上の例示に限定されない。以上の通り、目標音高Pを表現するために使用される有効音名データの個数Nは任意である。

[0068] D : 第4実施形態

第1、第2、および第3実施形態の生成部122は波形スペクトルを生成したが、第4実施形態では、生成部122が、生成モデルを用いて音信号Vを生成する。第4実施形態の機能構成は、図2と基本的に同じだが、合成部123は不要である。訓練部114は、参照信号Rを用いて生成モデルを訓練し、生成部122はその生成モデルを用いて音信号Vを生成する。第4実施形態における訓練用の発音単位データは、各発音単位の制御データXと参照信号Rの部分波形（すなわち参照信号Rのサンプル）のセットである。

[0069] 第4実施形態の訓練部114は、訓練データを受け取り、その各バッチの複数の発音単位の部分波形と制御データXとを順番に用いて生成モデルを訓練する。生成モデルは、サンプリング周期（時刻t）毎に、音信号Vのサンプルを表す出力データを推定する。訓練部114は、制御データXから推定された出力データの時系列と訓練データのうち対応する部分波形とに基づいて損失関数L（1バッチ分の累算値）を計算し、その損失関数Lが最小化されるように生成モデルの複数の変数を最適化する。こうして確立された生成モデルは、複数の発音単位データにおける各制御データXと、参照信号Rの部分波形との間に潜在する関係を学習している。

[0070] 第4実施形態の生成部122は、確立された生成モデルを用いて、制御データXに応じた音信号Vを生成する。生成部122は、生成モデルを用いて、サンプリング周期ごと（時刻t）に、制御データXに応じた音信号Vのサンプルを示す出力データを推定する。出力データが複数のサンプルの各々の確率密度分布を表す場合は、生成部122は、その成分の確率密度分布に従う乱数を生成して、当該乱数を音信号Vのサンプルとして出力する。出力データがサンプルの値を表す場合には、そのサンプルの時系列が音信号Vとして出力される。

[0071] E：第5実施形態

図2に図示される実施形態の音生成機能では、楽譜データの一連の発音単位の情報に基づいて、音信号Vを生成していたが、鍵盤等から供給される発音単位の情報に基づいて、リアルタイムに音信号Vを生成するようにしても

よい。その場合、処理部 121 は、各時点の制御データ X を、その時点までに供給された発音単位の情報に基づいて生成する。ここでは、制御データ X に含まれるコンテキストデータ X3 には、基本的に、未来の発音単位の情報を含むことができないが、過去の情報から未来の発音単位の情報予測して、未来の発音単位の情報を含めるようにしてもよい。

[0072] F：変形例

第 1 実施形態では、音高データ X1 を構成する M 個の音名データのうち、目標音高 P に近い基準音高 Q に対応する音名の音名データを有効音名データとして選択したが、目標音高 P から遠い基準音高 Q に対応する音名の音名データを有効音名データとして選択してもよい。その場合、第 1 実施形態のずれ値は、±50 セントを超える音高差を表現できるようにスケールされる。

[0073] 第 2 および第 3 実施形態では、目標音高 P と基準音高 Q との近接度は、セントスケールにおける両者間の音高差に応じて 0 から 1 までの範囲内においてリニアに変化した。代わりに、図 10 のような正規分布などの確率分布またはコサインカーブ等の任意のカーブ、または、折れ線に従って 1 から 0 に減少してもよい。

[0074] 第 1 および第 2 実施形態では、セントスケールで音名の対応付けを行っていたが、ヘルツスケールなどのその他音高を表現する任意のスケールで対応付けを行っても良い。その際、ずれ値は各スケール上における適切な値を挿入する。

[0075] 第 1、第 2 実施形態では、処理部 113 または処理部 121 において 0 から 1 にスケールしていたが、スケールについては任意の値で行っても良い。例えば、-1 から +1 にスケールしてもよい。

[0076] なお、音信号合成システム 100 が合成する音信号は、楽器音または音声に限らず、動物の鳴き声、または風音および波音のような自然界の音であっても、その音高の動的な制御を行いたい場合は、本開示を適用できる。

符号の説明

[0077] 1 0 0…音信号生成装置、1 1…制御装置、1 2…記憶装置、1 3…表示装置、1 4…入力装置、1 5…放音装置、1 1 1…解析部、1 1 2…時間合わせ部、1 1 3…処理部、1 1 4…訓練部、1 2 1…処理部、1 2 2…生成部、1 2 3…合成部。

請求の範囲

- [請求項1] 合成すべき第1音信号の音高を表す第1音高データを生成し、
 第2音信号の音高を表す第2音高データと前記第2音信号との関係
 を学習した生成モデルを用いて、前記第1音高データに応じた前記第
 1音信号を示す出力データを推定する
 コンピュータにより実現される音信号合成方法であって、
 前記第1音高データは、相異なる音名に対応する複数の音名データ
 を含み、
 前記第1音高データの生成においては、前記複数の音名データのう
 ち、前記第1音信号の音高に対応する音名データを、当該音名データ
 に対応する音名の基準音高と当該第1音信号の音高との相違に応じた
 ホット値に設定する
 音信号合成方法。
- [請求項2] 前記第1音高データの生成においては、前記複数の音名データのう
 ち、前記第1音信号の音高に対応する音名データ以外の音名データを
 、前記第1音信号の音高に関与しないことを表すコールド値に設定す
 る
 請求項1の音信号合成方法。
- [請求項3] 前記第1音信号の音高は動的に変化し、
 前記第1音高データは、前記第1音信号において動的に変化する音
 高を表す
 請求項1または請求項2の音信号合成方法。
- [請求項4] 前記第2音信号の音高は動的に変化し、
 前記第2音高データは、前記第2音信号において動的に変化する音
 高を表し、
 前記生成モデルは、前記第2音信号と前記第2音高データとを用い
 て訓練されている
 請求項1から請求項3の何れかの音信号合成方法。

- [請求項5] 前記第1音信号の音高は、一の音名に対応する発音期間内において動的に変化し、前記第1音信号の音高に対応する音名データに設定されるホット値は、当該音高に応じて変化する
- 請求項1から請求項3の何れかの音信号合成方法。
- [請求項6] 前記第1音信号の音高に対応する音名データは、相異なる音名に対応する複数の単位範囲のうち、当該音高を含む1個の単位範囲に対応する音名の音名データである
- 請求項1から請求項5の何れかの音信号合成方法。
- [請求項7] 前記第1音信号の音高に対応する音名データは、相異なる音名に対応する複数の基準音高のうち、当該音高を挟む2個の基準音高にそれぞれ対応する2個の音名データである
- 請求項1から請求項5の何れかの音信号合成方法。
- [請求項8] 前記第1音信号の音高に対応する音名データは、相異なる音名に対応する複数の基準音高のうち、当該音高に近いN個（Nは1以上の自然数）の基準音高にそれぞれ対応するN個の音名データである
- 請求項1から請求項5の何れかの音信号合成方法。
- [請求項9] 前記推定される出力データは、前記第1音信号の波形スペクトルに関する特徴量を表す
- 請求項1から請求項8の何れかの音信号合成方法。
- [請求項10] 前記推定される出力データは、前記第1音信号のサンプルを表す
- 請求項1から請求項8の何れかの音信号合成方法。
- [請求項11] 合成すべき音信号の音高を表す音高データを用意し、
- 前記音高データの入力に対して前記音信号を示す出力データが出力されるように生成モデルを訓練する
- コンピュータにより実現される生成モデルの訓練方法であって、
- 前記音高データは、相異なる音名に対応する複数の音名データを含み、
- 前記音高データの使用においては、前記複数の音名データのうち、

前記音信号の音高に対応する音名データを、当該音名データに対応する音名の基準音高と当該音信号の音高との相違に応じたホット値に設定する

訓練方法。

[請求項12] 1以上のプロセッサと1以上のメモリとを具備する音信号合成システムであって、

前記1以上のメモリは、第2音信号の音高を表す第2音高データと前記第2音信号との関係を学習した生成モデルを記憶し、

前記1以上のプロセッサは、

合成すべき第1音信号の音高を表す第1音高データを生成し、

前記生成モデルを用いて、前記第1音高データに応じた前記第1音信号を示す出力データを推定し、

前記第1音高データは、相異なる音名に対応する複数の音名データを含み、

前記1以上のプロセッサは、前記第1音高データの生成において、前記複数の音名データのうち、前記第1音信号の音高に対応する音名データを、当該音名データに対応する音名の基準音高と当該第1音信号の音高との相違に応じたホット値に設定する

音信号合成システム。

[請求項13] 合成すべき第1音信号の音高を表す第1音高データを生成する処理部、および、

第2音信号の音高を表す第2音高データと前記第2音信号との関係を学習した生成モデルを用いて、前記第1音高データに応じた前記第1音信号を示す出力データを推定する生成部

としてコンピュータを機能させるプログラムであって、

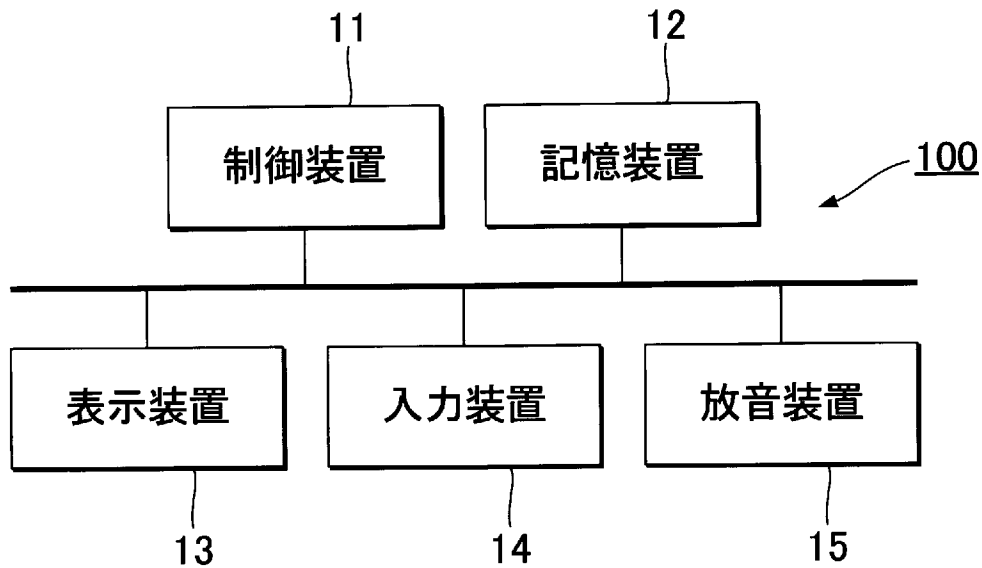
前記第1音高データは、相異なる音名に対応する複数の音名データを含み、

前記第1音高データの生成においては、前記複数の音名データのう

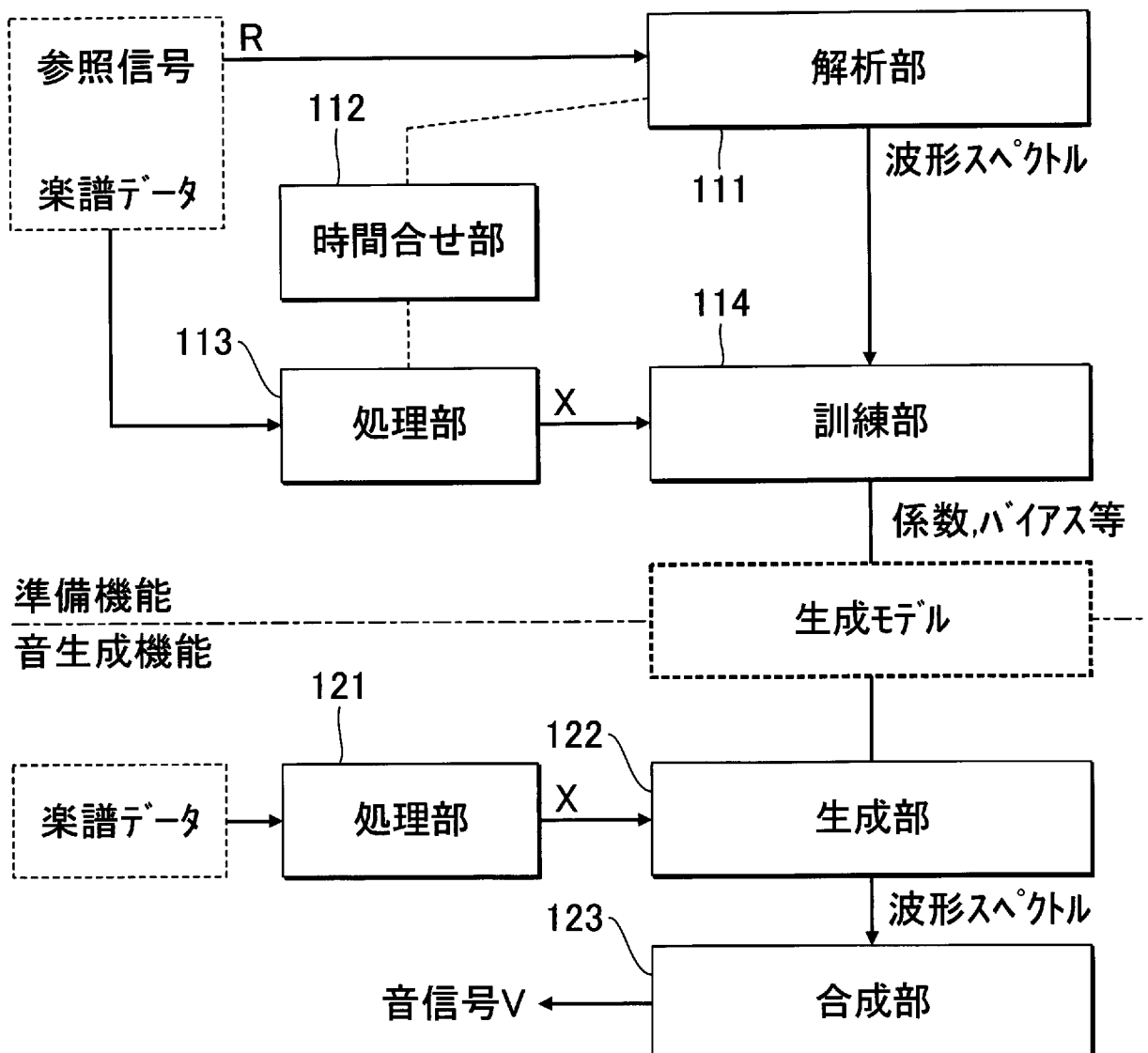
ち、前記第 1 音信号の音高に対応する音名データを、当該音名データに対応する音名の基準音高と当該第 1 音信号の音高との相違に応じたホット値に設定する

プログラム。

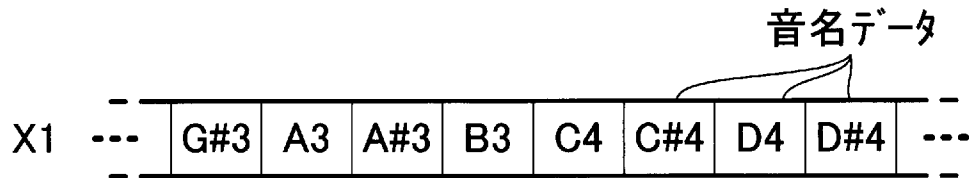
[図1]



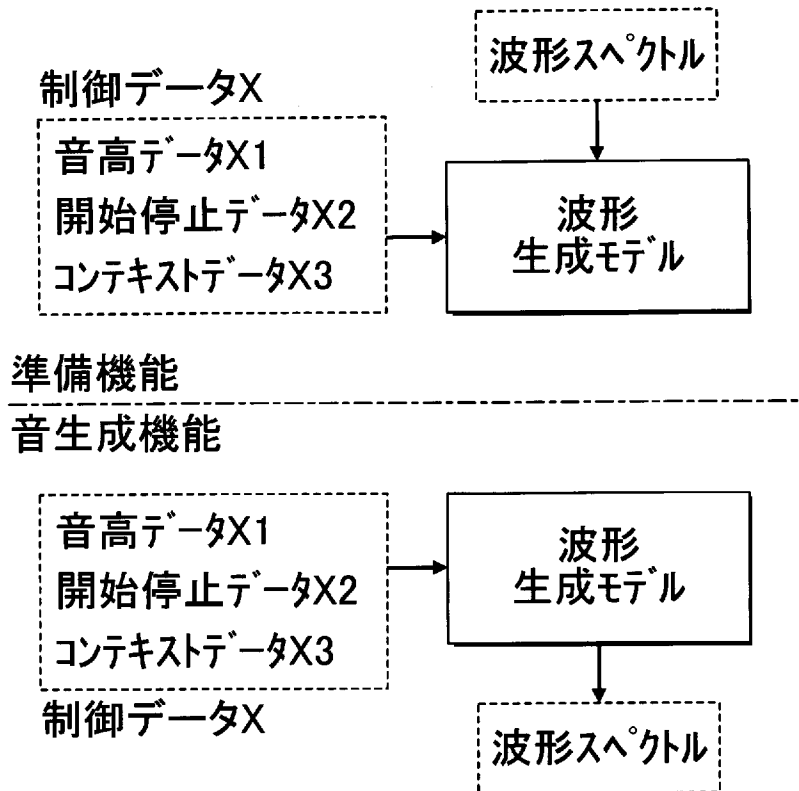
[図2]



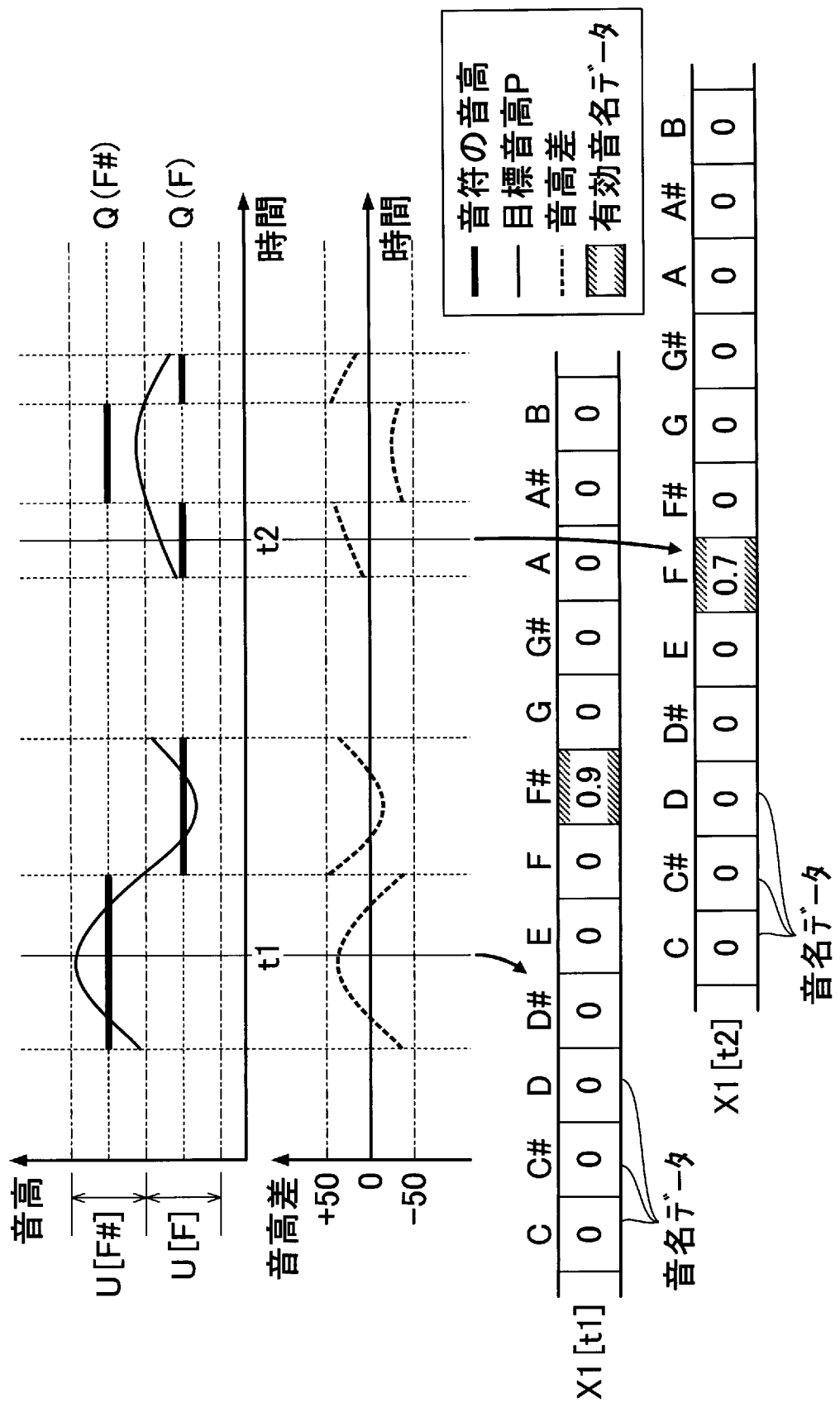
[図3]



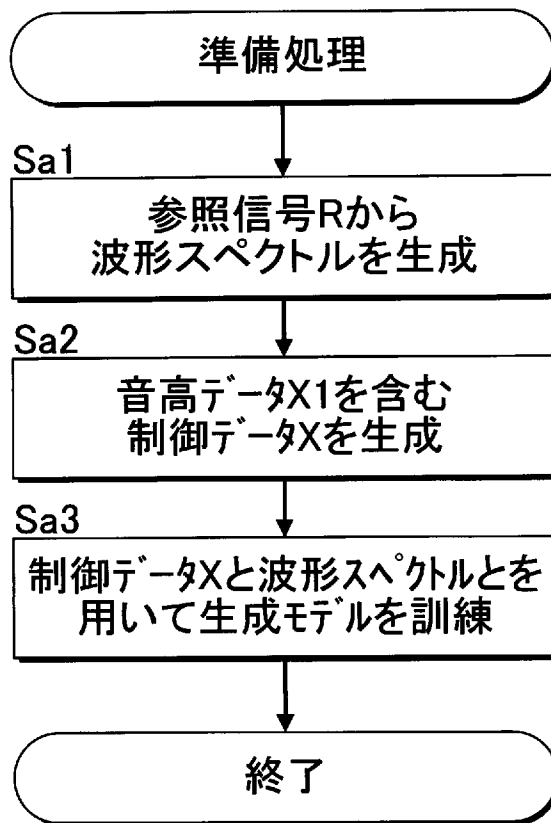
[図4]



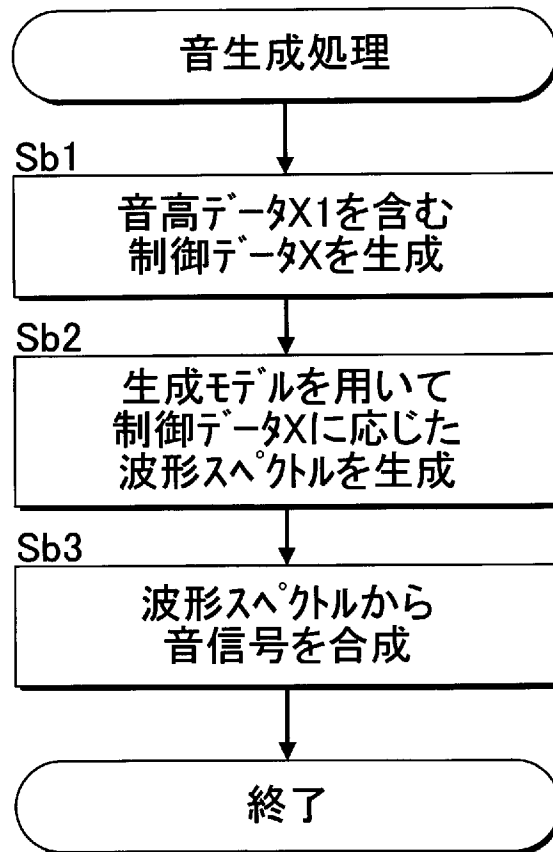
[図5]



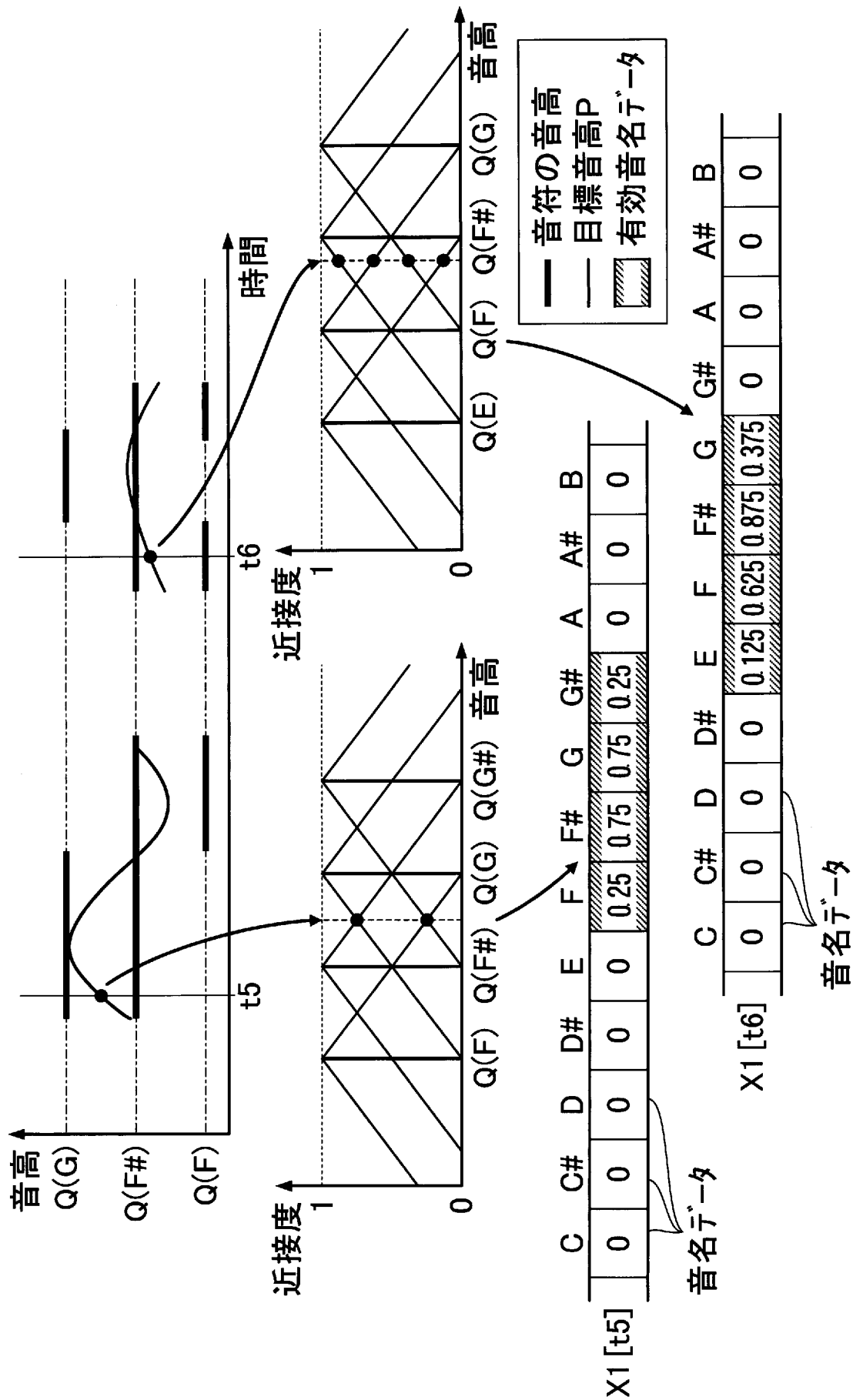
[図6]



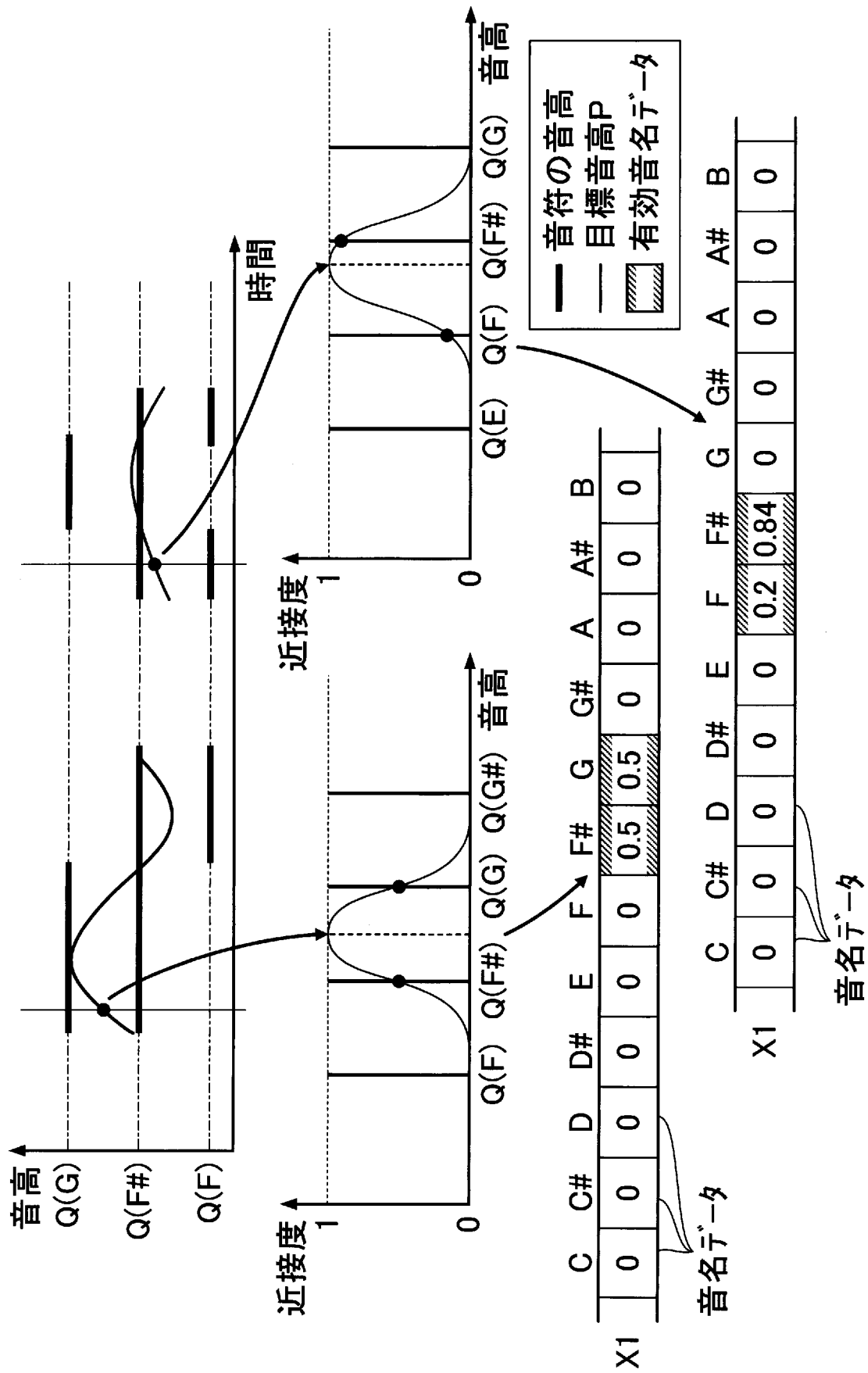
[図7]



[図9]



[図10]



INTERNATIONAL SEARCH REPORT

International application No.

PCT/JP2020/006162

A. CLASSIFICATION OF SUBJECT MATTER

G10L 13/00 (2006.01)i; G10H 7/08 (2006.01)i
FI: G10H7/08; G10L13/00 100Y

According to International Patent Classification (IPC) or to both national classification and IPC

B. FIELDS SEARCHED

Minimum documentation searched (classification system followed by classification symbols)

G10L13/00; G10H7/08

Documentation searched other than minimum documentation to the extent that such documents are included in the fields searched

Published examined utility model applications of Japan	1922-1996
Published unexamined utility model applications of Japan	1971-2020
Registered utility model specifications of Japan	1996-2020
Published registered utility model applications of Japan	1994-2020

Electronic data base consulted during the international search (name of data base and, where practicable, search terms used)

C. DOCUMENTS CONSIDERED TO BE RELEVANT

Category*	Citation of document, with indication, where appropriate, of the relevant passages	Relevant to claim No.
A	"How to use LSTM neural network with 3D categorical input and output?", STACK OVERFLOW, [online], 22 May 2018, [retrieved on 12 May 2020], <URL:https://stackoverflow.com/questions/50474524/how-to-use-lstm-neural-network-with-3d-categorical-input-and-output?rq=1>, in particular, see "Answer"	1-13
A	BLAAUW, Merlijn et al., "A Neural Parametric Singing Synthesizer Modeling Timbre and Expression from Natural Songs", Applied Sciences, [online], December 2017, 1313, pp. 1-23, [retrieved on 12 May 2020], <URL:https://pdfs.semanticscholar.org/0376/966eldee6ead97730f2f8ce55df936a7111a.pdf>, pp. 2-9, 19-20	1-13



Further documents are listed in the continuation of Box C.



See patent family annex.

* Special categories of cited documents:

"A" document defining the general state of the art which is not considered to be of particular relevance
 "E" earlier application or patent but published on or after the international filing date
 "L" document which may throw doubts on priority claim(s) or which is cited to establish the publication date of another citation or other special reason (as specified)
 "O" document referring to an oral disclosure, use, exhibition or other means
 "P" document published prior to the international filing date but later than the priority date claimed

"I" later document published after the international filing date or priority date and not in conflict with the application but cited to understand the principle or theory underlying the invention
 "X" document of particular relevance; the claimed invention cannot be considered novel or cannot be considered to involve an inventive step when the document is taken alone
 "Y" document of particular relevance; the claimed invention cannot be considered to involve an inventive step when the document is combined with one or more other such documents, such combination being obvious to a person skilled in the art
 "&" document member of the same patent family

Date of the actual completion of the international search
12 May 2020 (12.05.2020)

Date of mailing of the international search report
26 May 2020 (26.05.2020)

Name and mailing address of the ISA/
Japan Patent Office
3-4-3, Kasumigaseki, Chiyoda-ku,
Tokyo 100-8915, Japan

Authorized officer

Telephone No.

INTERNATIONAL SEARCH REPORT

International application No.

PCT/JP2020/006162

C (Continuation). DOCUMENTS CONSIDERED TO BE RELEVANT

Category*	Citation of document, with indication, where appropriate, of the relevant passages	Relevant to claim No.
A	US 10068557 B1 (GOOGLE LLC) 04.09.2018 (2018-09-04) column 5, line 54 to column 10, line 21	1-13
A	西村 方成, 外 4 名, Deep Neural Network に基づく歌声合成の検討, 日本音響学会 2016 年春季研究発表会講演論文集, [CD-ROM], 24 February 2016, pp. 213-214, entire text, all drawings, (NISHIMURA, Masanari et al., "Investigation of singing voice synthesis based on deep neural networks", Lecture proceedings of 2016 spring research conference of the acoustical Society of Japan)	1-13

INTERNATIONAL SEARCH REPORT
Information on patent family members

International application No.

PCT/JP2020/006162

Patent Documents
referred in the
Report

Publication
Date

Patent Family

Publication
Date

US 10068557 B1

04 Sep. 2018

WO 2019/040132 A1
paragraphs [0067]-
[0105]

A. 発明の属する分野の分類（国際特許分類（IPC）） G10L 13/00(2006.01)i; G10H 7/08(2006.01)i FI: G10H7/08; G10L13/00 100Y														
B. 調査を行った分野														
調査を行った最小限資料（国際特許分類（IPC）） G10L13/00; G10H7/08														
最小限資料以外の資料で調査を行った分野に含まれるもの														
<table border="0"> <tr> <td>日本国実用新案公報</td> <td>1922 - 1996年</td> </tr> <tr> <td>日本国公開実用新案公報</td> <td>1971 - 2020年</td> </tr> <tr> <td>日本国実用新案登録公報</td> <td>1996 - 2020年</td> </tr> <tr> <td>日本国登録実用新案公報</td> <td>1994 - 2020年</td> </tr> </table>			日本国実用新案公報	1922 - 1996年	日本国公開実用新案公報	1971 - 2020年	日本国実用新案登録公報	1996 - 2020年	日本国登録実用新案公報	1994 - 2020年				
日本国実用新案公報	1922 - 1996年													
日本国公開実用新案公報	1971 - 2020年													
日本国実用新案登録公報	1996 - 2020年													
日本国登録実用新案公報	1994 - 2020年													
国際調査で使用した電子データベース（データベースの名称、調査に使用した用語）														
C. 関連すると認められる文献														
引用文献の カテゴリー*	引用文献名 及び一部の箇所が関連するときは、その関連する箇所の表示	関連する 請求項の番号												
A	How to use LSTM neural network with 3D categorical input and output?, STACK OVERFLOW, [online], 2018.05.22, [retrieved on 2020.05.12], <URL: https://stackoverflow.com/questions/50474524/how-to-use-lstm-neural-network-with-3d-categorical-input-and-output?rq=1> 特に、「Answer」を参照	1-13												
A	Merlijn Blaauw et al., A Neural Parametric Singing Synthesizer Modeling Timbre and Expression from Natural Songs, Applied Sciences, [online], 2017.12, 1313, pp. 1-23, [retrieved on 2020.05.12], <URL: https://pdfs.semanticscholar.org/0376/966eldee6ead97730f2f8ce55df936a7111a.pdf> 第2-9, 19-20ページ	1-13												
A	US 10068557 B1 (GOOGLE LLC) 04.09.2018 (2018 - 09 - 04) 第5欄第54行 - 第10欄第21行	1-13												
A	西村 方成, 外4名, Deep Neural Networkに基づく歌声合成の検討, 日本音響学会 2016年春季研究発表会講演論文集, [CD-ROM], 2016.02.24, 第213-214ページ 全文 全図	1-13												
☐ C欄の続きにも文献が列挙されている。 ☒ パテントファミリーに関する別紙を参照。														
<table border="0"> <tr> <td>* 引用文献のカテゴリー</td> <td>"T" 国際出願日又は優先日後に公表された文献であって出願と抵触するものではなく、発明の原理又は理論の理解のために引用するもの</td> </tr> <tr> <td>"A" 特に関連のある文献ではなく、一般的な技術水準を示すもの</td> <td>"X" 特に関連のある文献であって、当該文献のみで発明の新規性又は進歩性がないと考えられるもの</td> </tr> <tr> <td>"E" 国際出願日前の出願または特許であるが、国際出願日以後に公表されたもの</td> <td>"Y" 特に関連のある文献であって、当該文献と他の1以上の文献との、当業者にとって自明である組合せによって進歩性がないと考えられるもの</td> </tr> <tr> <td>"L" 優先権主張に疑義を提起する文献又は他の文献の発行日若しくは他の特別な理由を確立するために引用する文献（理由を付す）</td> <td>"&" 同一パテントファミリー文献</td> </tr> <tr> <td>"O" 口頭による開示、使用、展示等に言及する文献</td> <td></td> </tr> <tr> <td>"P" 国際出願日前で、かつ優先権の主張の基礎となる出願の日の後に公表された文献</td> <td></td> </tr> </table>			* 引用文献のカテゴリー	"T" 国際出願日又は優先日後に公表された文献であって出願と抵触するものではなく、発明の原理又は理論の理解のために引用するもの	"A" 特に関連のある文献ではなく、一般的な技術水準を示すもの	"X" 特に関連のある文献であって、当該文献のみで発明の新規性又は進歩性がないと考えられるもの	"E" 国際出願日前の出願または特許であるが、国際出願日以後に公表されたもの	"Y" 特に関連のある文献であって、当該文献と他の1以上の文献との、当業者にとって自明である組合せによって進歩性がないと考えられるもの	"L" 優先権主張に疑義を提起する文献又は他の文献の発行日若しくは他の特別な理由を確立するために引用する文献（理由を付す）	"&" 同一パテントファミリー文献	"O" 口頭による開示、使用、展示等に言及する文献		"P" 国際出願日前で、かつ優先権の主張の基礎となる出願の日の後に公表された文献	
* 引用文献のカテゴリー	"T" 国際出願日又は優先日後に公表された文献であって出願と抵触するものではなく、発明の原理又は理論の理解のために引用するもの													
"A" 特に関連のある文献ではなく、一般的な技術水準を示すもの	"X" 特に関連のある文献であって、当該文献のみで発明の新規性又は進歩性がないと考えられるもの													
"E" 国際出願日前の出願または特許であるが、国際出願日以後に公表されたもの	"Y" 特に関連のある文献であって、当該文献と他の1以上の文献との、当業者にとって自明である組合せによって進歩性がないと考えられるもの													
"L" 優先権主張に疑義を提起する文献又は他の文献の発行日若しくは他の特別な理由を確立するために引用する文献（理由を付す）	"&" 同一パテントファミリー文献													
"O" 口頭による開示、使用、展示等に言及する文献														
"P" 国際出願日前で、かつ優先権の主張の基礎となる出願の日の後に公表された文献														
国際調査を完了した日 12.05.2020	国際調査報告の発送日 26.05.2020													
名称及びあて先 日本国特許庁(ISA/JP) 〒100-8915 日本国 東京都千代田区霞が関三丁目4番3号	権限のある職員（特許庁審査官） 岩田 淳 5Z 4052 電話番号 03-3581-1101 内線 3591													

PCT/JP2020/006162

様式 PCT/ISA/210 (パテントファミリー用別紙) (2015年1月)