

音声のカオス・フラクタル性とその生成モデル

古賀 博之*・中川 匡弘*

Chaotic and Fractal Properties in Vocal Sound and its Synthesis Model

Hiroyuki KOGA*, Masahiro NAKAGAWA*

In the field of synthesis of vocal sounds, various models composed of glottal sound source, vocal tract, and radiation into space, have been proposed. These models, however, do not consider how to involve the chaos property into them, which are not chaotic. In the present work, we proposed an oscillation model of voice synthesizer with chaotic behavior. This model utilized an exponential spring and damping in the self-oscillating model of the vocal cords. The present results show that the synthesized voices involve the chaotic property by means of the Lyapunov analysis and the analysis of surrogate data.

Key Words: Vocal sounds, Chaos, Lyapunov dimension, Surrogated data

1. 序 論

カオスの概念¹⁾には三つの特徴がある。それらは軌道不安定性 (orbital instability), 長期予測不能性 (long-term unpredictability), そして自己相似性 (self-similarity) であり, これらの特徴を表現するための様々な研究が進められている。軌道不安定性と長期予測不能性は初期値に対する鋭敏性であり, リアプノフ指数 (Lyapunov exponents) およびKSエントロピー (Kolmogorov-Sinai entropy) により評価できる。また, アトラクタの自己相似構造である自己相似性は, 非整数のフラクタル次元²⁾という尺度で特徴づけられる。

自然界に存在する様々な事象は, フラクタル性を有している。人間の発する音声波形にもゆらぎが存在しており, 自己相似性があるためフラクタル的な振舞いがあると指摘されている³⁾。音声や脳波などの時系列データのリアプノフ指数を算出する方法が, 佐野ら⁴⁾, Eckmannら⁵⁾, およびWolfら⁶⁾によって独立に提案されている。この音声のフラクタル性は, リアプノフスペクトラム解析により定量化できる⁷⁾。音声の解析結果から, 最大リアプノフ指数が正となるため, 音声にはカオスが内在しているといわれている。

一方, 音声に関する研究においては, その発生機構のモデル化や音声の特徴抽出および認識, 音声情報処理に関する手法が精力的に進められている。音声生成の分野においては, 機械振動モデルを基にした音声を人

工的に生成する手法として, 様々な音声生成モデル⁸⁻¹⁰⁾の研究が行われている。しかしながら, これらのモデルの生成波形は, 周期的であり, 本来の音声を持つゆらぎをもっていない。人間にとって周期的な音は, 機械的な不自然な感覚および不快感をもたらす¹¹⁾。

これまで, 我々の研究室では, カオス的なゆらぎをもつような音声を生成するモデルに関する研究を行ってきた。特に, 山口ら¹²⁾によって声帯の1質量モデルを非線形振動させることで, カオス的なゆらぎをもった音源を実現するモデルが提案されている。このモデルでは, 1周期のみを切り出し連続させた完全に周期的な音声波形と, 生成されたカオスゆらぎを合成し, 自然な音声生成を行っている。

今回, 我々は石坂の研究による2質量モデル¹⁰⁾に対して, 山口のモデルを改良した指数関数型の非線形項を導入し, カオスの振舞いを示す音声生成モデル^{13,14)}を構築した。さらに生成音声のリアプノフ解析およびサロゲーション法 (the method of surrogated data)^{15,17)}を組み合わせて, 生成音声のカオスであることを検証した。

本論文では, 次節で, 石坂による声帯の2質量モデルとその運動方程式を示し, 次いで, 第3節では, 我々が今回導入したバネと粘性に関する指数関数型の非線形項について述べる。第4節では, 肺, 声門 (声帯), 声道, 口などの発声器官を電氣的等価回路に置き換えたモデルおよび, 声門内のインピーダンスと声帯モデルの関係を示す。第5節では, シミュレーション方法を説明し, 第6節でその結果とカオス性の検証を行う。最終節では, 以上のまとめと課題について述べる。

原稿受付: 平成11年5月21日

*長岡技術科学大学電気系

2. 声帯振動モデル

我々は、基本的なFlanaganらの1質量モデル⁸⁾に比べてより生理学的な振舞いが可能な石坂らの2質量モデル¹⁰⁾を用いる。このモデルの構成をFig. 1に示す。声帯の形状は左右対称形と仮定し、ここでは片方の運動のみを議論する。

声門は器官と声道の間に位置し、声帯の振動によって形成される隙間である。また、入口や出口での収縮および膨張の規模は、声帯の変位に依存して決定される。

Fig. 1のモデルにおいて、2つの声帯は各々質量 m_1 , m_2 、バネの復元力 ψ_1 , ψ_2 、粘性 ξ_1 , ξ_2 をもつ簡単な機械振動子として構成され、各々の質量の変位を x_1 , x_2 とする。また、隣り合う声帯 m_1 と m_2 は線形な連結バネ s_c で接続される。さらに、Fig. 1のように、肺で発生する声門下圧を P_s 、声帯の有効長を l_g 、太さを d_1 , d_2 、声門の隙間の横断断面積を A_{g1} , A_{g2} 、そして声門を通過する体積流の平均を U_g と表す。

変位 x_1 , x_2 を用いると、向かい合う声帯どうしの間隔 h_1 , h_2 は、各々、 $h_i = x_i + A_{g0}/2$, $l_g = x_i - x_j$, ($i=1,2$)と表現できる。ここで x_c は声帯が閉じる（衝突が起こる）ときの質量の変位であり、 A_{g0} は自然な発声（phonation neutral）時の横断断面積を表す。従って、横断断面積 A_{g1} , A_{g2} は $A_{gi} = 2 l_g h_i$, ($i=1,2$)と表現できる。

各々の声帯表面に圧力 P_1 , P_2 が加わるとすると、各質量の運動方程式は、

$$\frac{d^2 x_i}{dt^2} + 2 \omega_i \xi_i(x_i) \frac{dx_i}{dt} - \omega_i^2 \psi_i(x_i) = \frac{F_i - s_c(x_i, x_j)}{m_i}, \quad (1)$$

$$\omega_i^2 = \frac{k_i}{m_i}, \quad F_i = P_i l_{gi}, \quad (i, j=1, 2, i \neq j),$$

で表現される。式(1)で、 ω_i は各声帯の振動周波数、 F_i は各声帯を押し上げる強制力、そして k_i は線形バネ定数である。実際には、連結バネの力の作用する方向は斜めであるが、本研究では、我々は質量の深さを無視して、簡単の為に $s_c(x_1, x_2) = k_c(x_1 - x_2)$ を用いる。ここで k_c は連結バネの線形バネ定数である。

3. 非線形バネと非線形粘性

カオス音声生成モデルは、声帯を非線形振動させてカオス的なゆらぎをもつ音声を生成することを目的としている。声帯振動を非線形にするためには、式(1)の運動方程式の振動項（バネの復元力 $\psi_i(x_i)$ および粘性 $\xi_i(x_i)$ ）を非線形項として定義する必要がある。

非線形関数としていくつかの形が考えられる。一般的に変位 x のべき乗を利用した多項式表現では、零近傍では希望する非線形性が保たれるが x の変位が大きくなるとだんだんと異なったより高次の非線形性に近づいてしまうという欠点がある。そこで我々は、非線形項として指数関数を用いる。指数関数は前述のような欠点はなく、変位が大きい場合でも非線形性は保持され、さらに微分演算が容易であるという長所を併せ持つ。

最初に非線形バネの復元力 ψ について述べる。このバネの定義には、戸田格子の指数関数型ポテンシャルの概念¹⁸⁾を利用する。即ち、非線形バネのポテンシャル $\phi(x)$ は、定数 A , a , b , c (a , b の符号は同じと仮定する) に対して、 $\phi(x) = A \exp(-bx) + ax + c$ と表現される。とりわけ、4つの定数を $A = a/b$, and $\varepsilon_1 = (-a) = (-1/b) > 0$ と仮定すると、声帯の変位 x_i ($i=1,2$) におけるバネの復元力は、これらのポテンシャルを x_i に関して微分し、即ち、

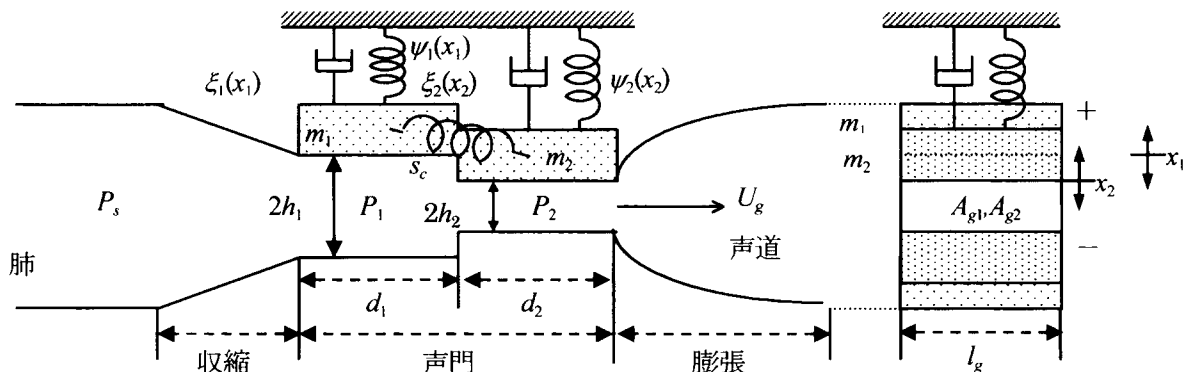


Fig. 1 Schematic diagram of the two-mass approximation of the vocal cords.

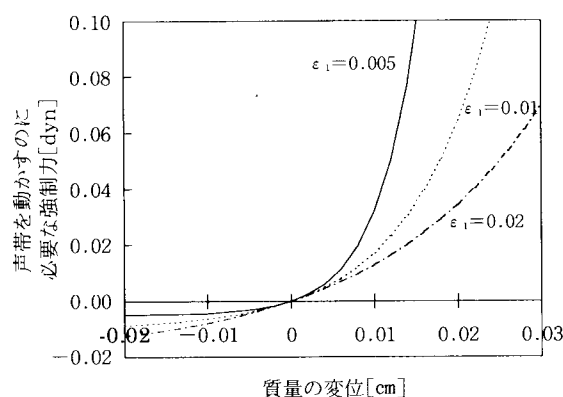


Fig. 2 Characteristics of the nonlinear stiffness.

$$\psi(x_i) = -\frac{d\phi(x_i)}{dx_i} = \varepsilon_1 \left\{ 1 - \exp\left(-\frac{x_i}{\varepsilon_1}\right) \right\}, \quad (2)$$

と定義される。Fig. 2に、この非線形バネを動かすのに必要な強制力 ($-\psi(x_i)$) の変化のグラフを示す。同図では、声帯を拡げる ($x_i > 0$) には大きな力が必要で、逆に声帯を閉じる ($x_i < 0$) には弱い力で十分なことを示している。また、パラメータ ε_1 は非線形性を制御する役割をもつ。

次に、非線形粘性 ξ の定義式について述べる。一般的な振動系では粘性率は一定（線形）である。本モデルも根本的には同様の考えを基本とするが、向かい合う声帯どうしが衝突する際の反発力の生成に、粘性の定義式を利用する。すなわち、声帯どうしが衝突する際 ($x_i < x_c$, $x_c = -A_{g0}/2l_g$) に結果的に運動方程式に反発力（または粘着力）を与えるような構造である。従って、非線形粘性は、一定の粘性率 ζ_i と ε_2 をパラメータとする指数関数によって次のように定義される。

$$\xi(x_i) = \zeta_i + \exp\left\{-\frac{x_i - x_c}{\varepsilon_2}\right\}, \quad (3)$$

ここで、 ε_2 は正の定数である。この式(3)による粘性表現をFig.3に示す。同図において、粘性率の変化は

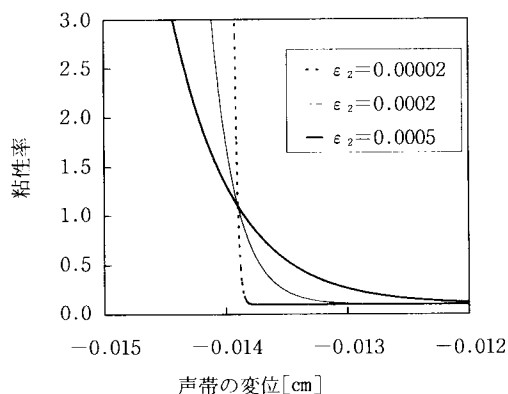


Fig. 3 Characteristics of the exponential damping.

声帯が開く方向では定数に収束し、衝突発生時 ($x_i = x_c$) に臨界減衰 (critical damping) となる。さらに閉鎖時には結果的に強い反発力となるように、パラメータ ε_2 の制御によって粘性率は大きく変化する。

4. 電気的等価回路と声帯圧力

肺から口における発声（放射）までの人間の発声器官は、Fig. 4に示されるような電気的等価回路と其中に存在する様々な音響インピーダンス成分によって表現できる。声門内では、体積流 U_g が一様であると仮定すれば、声門内部のインピーダンス成分は、空気密度を ρ 、空気の動粘性係数を μ とすると石坂のモデル¹⁰⁾と同様に、

$$\begin{aligned} R_c &= 1.37R_b |U_g| h_1^{-2}, R_{v1} = R_v d h_1^{-3}, \\ R_e &= -0.5R_b |U_g| h_2^{-2}, L_{g1} = L_g d h_1^{-1}, \\ R_{12} &= R_b |U_g| (h_2^{-2} - h_1^{-2}), (i=1,2), \end{aligned} \quad (4)$$

と表現できる。ここで、各定数は、 $R_b = \rho / 8l_g^2$, $R_v = 3\mu / 2l_g$, $L_g = \rho / 2l_g$ である。式(4)で R_c は声門入口での収縮、 R_{v1} および R_{12} は声帯 m_1 による圧力減衰、 R_{12} は2つの声帯間で発生する速度の急速的な変化に

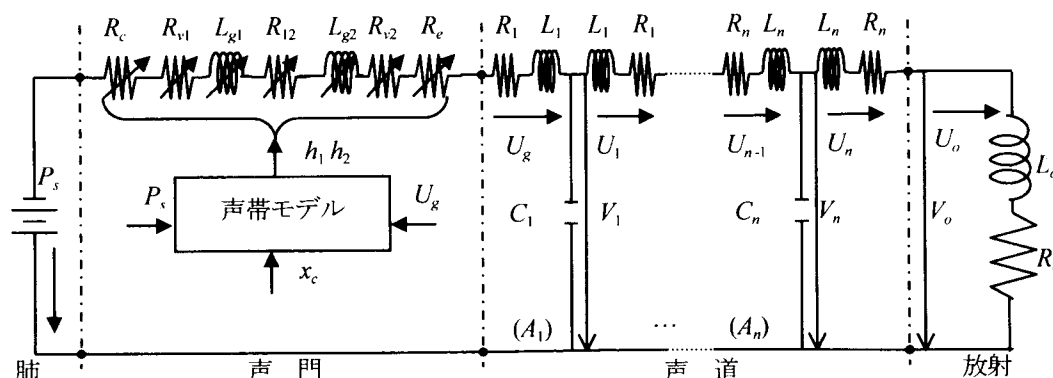


Fig. 4 Network model for the synthesis of voiced sounds.

対応し、 R_{v2} 、 L_{g2} は声帯 m_2 による圧力減衰、 R_e は声門出口での膨張を表現する。

声道内の各インピーダンスと放射インピーダンスは、声道の形状に依存する。すなわち、人間が声道と呼ばれる調音器官の形状を変化させて音声と作り出すのと同様に、声道や口のモデルも声道の形状に依存してインピーダンスを変化させ、実際の発声に近い機構を構成する。声道の長さを n 個所に分離したそれぞれの断面における横断面積を $A_1 \cdots A_n$ とし、その長さを l_n とすると、各インピーダンスは、

$$\begin{aligned} R_j &= S_j A_j^2 \sqrt{\rho \mu \omega_1 / 2}, L_j = 0.5 l_j \rho A_j^{-1}, \\ C_j &= l_j d_j / \rho c^2, (j=1, \cdots, n), \\ R_o &= \rho \omega_1^2 / 2c\pi, L_o = 8 \rho / 3\pi \sqrt{\pi A_n} \end{aligned} \quad (5)$$

と表現できる。ここで、 c は音速を表し、 S_j は j 番目の声道断面の円周長である。この等価回路モデルにおいて、放射インピーダンス R_o および L_o に加わる音圧波形 V_o が音声波形に対応するので、生成音声波形として取り扱う。

続いて、声帯質量の表面を押し上げる力となる圧力($h_1, h_2 > 0$)について考える。Fig. 4の電気的等価回路をみると、声帯 m_1 に加わる圧力 P_1 は肺から声帯 m_1 (R_{v1} と L_{g1} による平均)までの圧力分布を等価回路から求めればよく、 P_2 も同様である。

$$\begin{aligned} P_1 &= P_s - \left(R_e + \frac{R_{v1}}{2} \right) U_g - \frac{L_{g1}}{2} \cdot \frac{dU_g}{dt}, \\ P_2 &= P_1 - \left\{ R_{l2} + \frac{(R_{v1} + R_{v2})}{2} \right\} U_g - \frac{L_{g1} + L_{g2}}{2} \cdot \frac{dU_g}{dt}. \end{aligned} \quad (6)$$

次に声門が閉じている場合、圧力分布を厳密に求めるのは困難である。しかしながら、 R_e が無視できるものと仮定すると、近似的に声門加圧が閉鎖表面に対して直接作用すると考えられる。すなわち、 $h_1 \leq 0$ のときは $P_1 = P_s$ 、 $P_2 = 0$ と近似し、 $h_1 > 0$ かつ $h_2 \leq 0$ のときは $P_1, P_2 = P_s$ と近似する。

5. シミュレーション

我々はルンゲクッタ法により本モデルによる音声生成をシミュレートした。本節では、シミュレーションの流れの順に各々の微分方程式を示し説明する。以下の説明ではルンゲクッタ法の刻み時間を Δt と表現する。

1) 初期設定：声道の形状を決定し、それに伴い声道

および放射インピーダンスを式(5)により設定する。また、各モデル内の音圧、体積流、声帯変位などについて初期設定を行う。

2) 現在の声門間隔 h_1, h_2 から声門内部の各インピーダンスの値を計算する。ただし、声門が閉鎖(衝突)している場合、各インピーダンスは計算不可能であるため次の3)では式(7)は使用できず、かわりに式(9)を用いる。

3) 声門の体積流 U_g ： U_g の微分値を次式から求める。

$$\frac{dU_g}{dt} = \frac{P_s - V_1 - (R_e + R_{v1} + R_{l2} + R_{v2} + R_e + R_l)U_g}{L_{g1} + L_{g2} + L_l}. \quad (7)$$

ただし、声門が閉じている場合は、体積流の流れを停止させるように、微分の定義式、

$$\frac{dU_g(t)}{dt} = \lim_{\Delta t \rightarrow 0} \frac{U_g(t + \Delta t) - U_g(t)}{\Delta t} \quad (8)$$

に対して次の時刻での体積流 $U_g(t + \Delta t)$ を零とおいて得られる次式を用いる。

$$\frac{dU_g}{dt} = - \frac{U_g}{\Delta t} \Big|_{\Delta t \rightarrow 0} \quad (9)$$

4) 声道内の体積流 U_j ：

$$\begin{aligned} \frac{dU_j}{dt} &= \frac{V_j - V_{j+1} - (R_j + R_{j+1})U_j}{L_j + L_{j+1}}, \\ \frac{dU_n}{dt} &= \frac{V_n - (R_n + R_o)U_n}{L_j + L_o}, (j=1, \cdots, n-1). \end{aligned} \quad (10)$$

5) 声道内の音圧 V_j ：

$$\frac{dV_1}{dt} = \frac{U_g - U_1}{C_1}, \frac{dV_j}{dt} = \frac{U_{j-1} - U_j}{C_j}, (j=2, \cdots, n). \quad (11)$$

6) 放射音圧 V_o ：

$$V_o = R_o U_n + L_o \frac{dU_n}{dt} \quad (12)$$

7) 声帯の変位 x_i ：次の2階微分方程式から求められる。

$$\begin{aligned} \frac{d^2 x_i}{dt^2} &= \frac{F_i - s_c(x_i, x_{3-i})}{m_i} \\ &\quad - 2\omega_i \xi_i(x_i) \frac{dx_i}{dt} + \omega_i^2 \phi_i(x_i), (i=1, 2). \end{aligned} \quad (13)$$

8) 最後にルンゲクッタ法により以上の微分方程式の解を求め、2) へ戻り、続けて次の時刻での各要素の計算を実行する。

今回のシミュレーションで使用した定数をTable 1に示す。刻み時間 Δt は 10^{-6} (sec)とし、声道の横断断面積についてはFantによるX線写真による測定結果¹⁹⁾を利用した5種類の母音/a/, /i/, /u/, /e/, /o/について音声生成を実施した。生成音声データは、サンプリング周波数44(kHz)で再生し、生成開始後の安定したとみなせる区間を0.5(sec)後と考え0.5(sec)以降の波形を次節における音声のカオス・フラクタル解析に利用する。

6. シミュレーション結果とカオスの検証

母音/a/, /i/, /u/, /e/, /o/についての生成音声波形をFig. 5に示す。この図で生成音声波形はゆらぎを含んでいることがみられるが、視覚的にゆらぎの程度やそれがカオス的なゆらぎなのかということの判断は困難である。

そこで我々は、ゆらぎの存在やフラクタルの性質である自己相似性を視覚的に確認するためにアトラクタの埋め込みを行った。アトラクタは音声の時系列データ $f(t)$ を3次元空間に座標 $(x, y, z) = (f(t), f(t - \tau), f(t - 2\tau))$ と埋め込まれた連続する座標点である。

このとき、遅延時間 τ は時系列データの自己相関関数の最初に零となる点を利用して決定し²⁰⁾、次の3種類の音声信号に対しアトラクタを構成し形状の比較を実施した。

Table 1 Variables used in the present model.

	内 容	値	単 位
m_1	下側声帯の質量	0.125	g
m_2	上側声帯の質量	0.025	g
d_1	下側声帯の長さ	0.25	cm
d_2	上側声帯の長さ	0.05	cm
l_g	声 門 の 隙 間	1.4	cm
k_1	バネ 1 の線形係数	80000	dyn/cm
k_2	バネ 2 の線形係数	8000	dyn/cm
k_c	連結バネの係数	25000	dyn/cm
ε_1	バネの非線形係数	0.1	—
ε_2	粘性の非線形係数	0.000021	—
ρ	空 気 の 密 度	1.14×10^{-3}	g/cm ³
μ	空気の動粘性係数	1.86×10^{-4}	dyn s/cm ²
c	音 速	3.5×10^4	cm/s
ζ_1	粘性 1 の粘性係数	0.1	—
ζ_2	粘性 2 の粘性係数	0.6	—
P_s	声 門 下 圧	8	cmH ₂ O*

*1 (cmH₂O) = 980.7 (dyn/cm²)

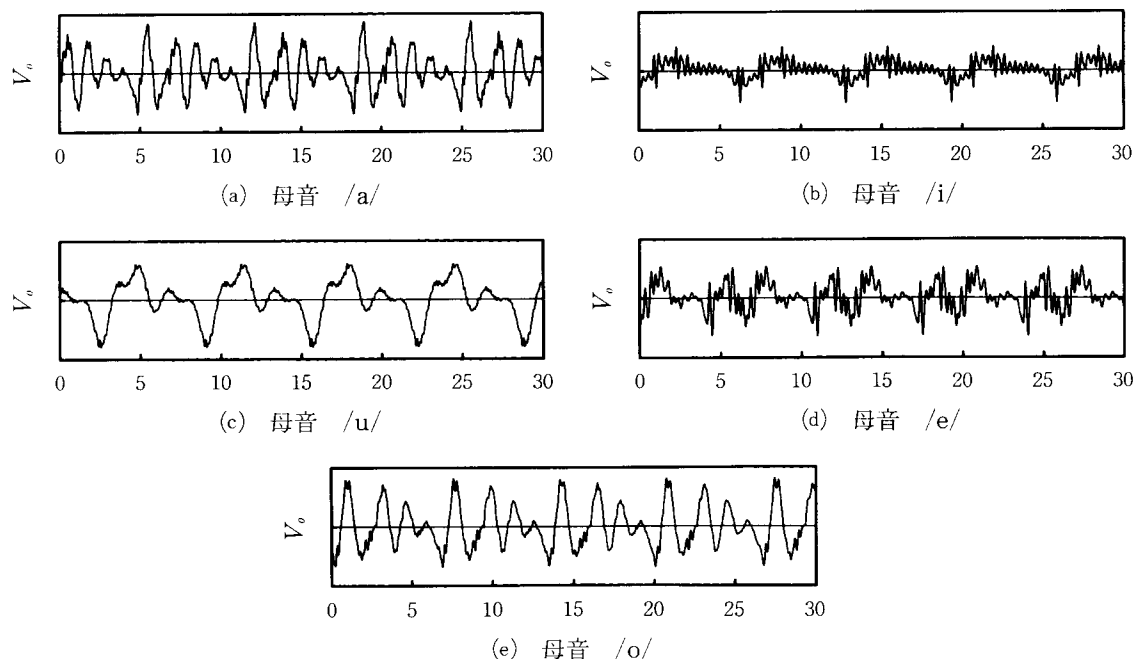


Fig. 5 The results of the simulation for vowels /a/, /i/, /u/, /e/, and /o/.

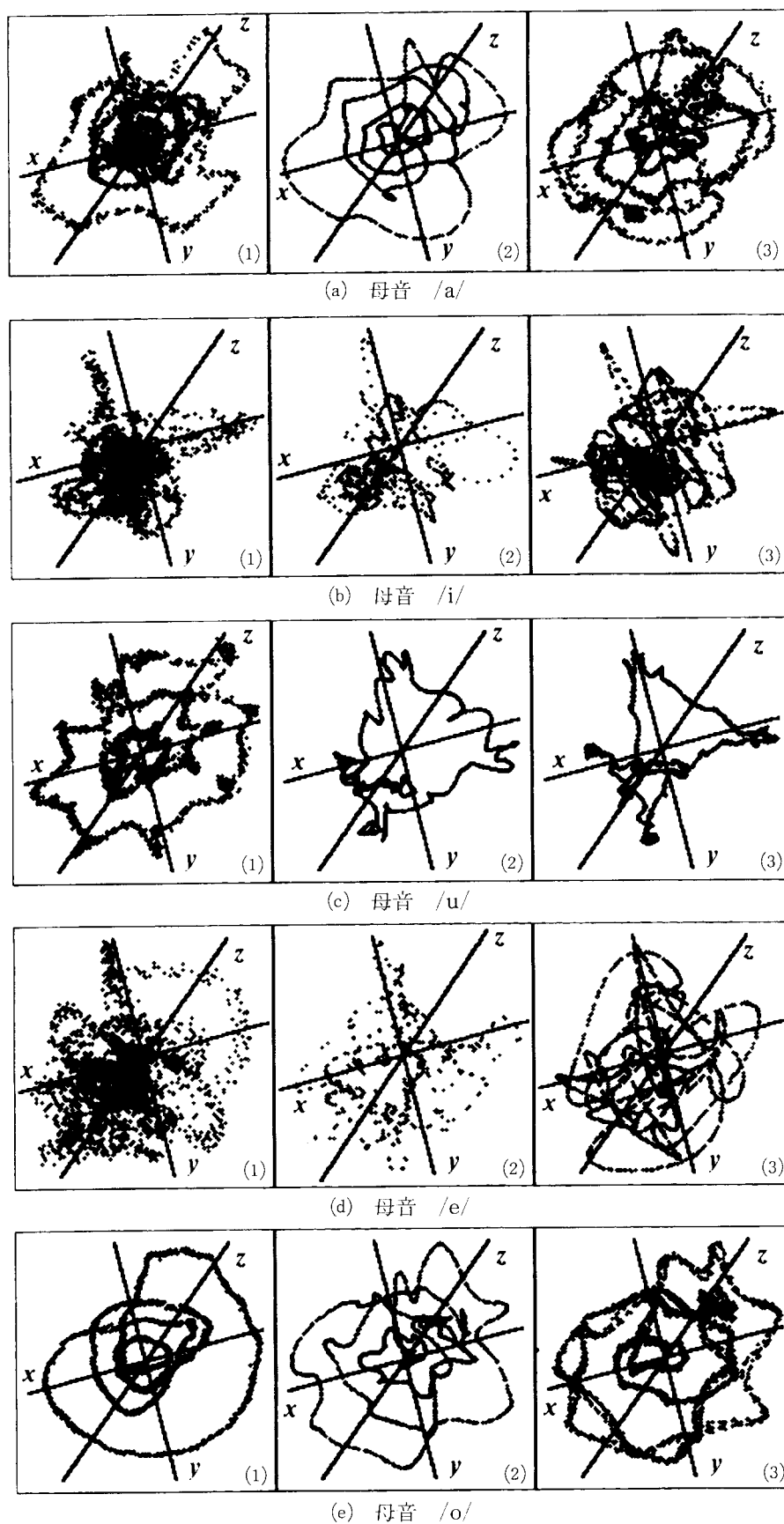


Fig. 6 The attractors of vocal sounds for vowels /a/, /i/, /u/, /e/, and /o/. This includes (1)human vocal sounds,(2)synthesized voices by Ishizaka's model, and (3)present synthesized voices.

- 1) 日本人成人男性の音声
- 2) 石坂の2質量モデル¹⁰⁾による生成音声
- 3) 本モデルによる生成音声

これらの音声の中で1)の音声は、細かい声道形状や声道自体の長さなどの要素がFantの測定値とは異なる為、定量的比較は困難であることに注意しなくてはならない。2)および3)において声道の形状はFantの測定値を利用して音声生成を行った。実音声はパーソナルコンピュータ(EPSON:PC-286VS)と音声入出力ボード(Canopus Corp.:AUDIOPAQ-98)を用いて、44(kHz)のサンプリング周波数、16(bits)の精度で採取した。2つのモデルによる生成音声は、定常状態であると判断できる、生成開始0.5(sec)以降の連続する3000点を埋め込んだ。

それぞれの母音についてのアトラクタをFig. 6に示す。ゆらぎの存在および自己相似性はアトラクタが帯状になっているかで視覚的に確認できる。実際の音声のアトラクタを観ると全ての母音ともに太い帯状である。この帯の太さがフラクタル性の自由度の高さを表現する。石坂モデルの生成音声は母音/i/, および、/e/については、点状に散布しているため、周期的な波形であることがわかる。しかしながら、石坂モデルの他のアトラクタは細いひも状であることから、基本周期ごとに位相がずれていると考えられる。これは、石坂によるモデルのバネの復元力が x^3 に比例しているためである。しかしながら、この非線形項の効果は小さく、時間と共に位相のずれは衰退する(周期波形に近づく)。

これに対して、カオス音声生成モデルの生成音声のアトラクタは、形状が石坂のモデルに類似しており、太い帯状となっていることから、位相および振幅のゆらぎを含んでいると視覚的に推定できる。さらに実際の音声と比較すると形状は多少異なるが、帯の太さは同程度であり、類似したフラクタル特性をもっていると考えられる。アトラクタの形状の異なる原因には、声道形状および長さ、アトラクタの視点位置、生成モデル本来の特性、などに由来する相違が考えられる。また特に、母音/i/および/e/についてのアトラクタが複雑な構造をしていることから、これらのカオスの振る舞いは強く、母音/u/や/o/に関しては比較的単純な構造なので弱いと予想される。

続いて、我々は連続する 10^4 (steps)の音声データに関してリアプノフ解析を行い、音声のカオス特性を定量的に評価した。カオス特性は、最大リアプノフ指数が正となる(軌道不安定性)、また、リアプノフ次元が非整数となる(自己相似性)ことから特徴づけられ

る。実際の音声データと本モデルであるカオス音声生成モデルの生成音声に対する最大リアプノフ指数の計算結果をTable 2に示す。その結果から、全ての算出値が正の値となるので、音声および生成音声の軌道不安定性を有しているといえる。さらに、生成音声の最大リアプノフ指数は、実際の音声の指数値に近い。これに対して、石坂のモデルのリアプノフ指数は全て零または負の値となったので、周期波形であるといえる。

Table 2 The maximum Lyapunov exponents of vowels.

母 音	音 声(bit)	生成音声(bit)
/a/	+0.032	+0.054
/i/	+0.156	+0.174
/u/	+0.030	+0.038
/e/	+0.084	+0.144
/o/	+0.013	+0.010

Table 3 The Lyapunov Dimensions of vowels.

母 音	音 声(bit)	生成音声(bit)
/a/	1.38	1.46
/i/	2.19	2.41
/u/	1.14	1.33
/e/	2.03	2.34
/o/	1.09	1.06

Table 4 The maximum Lyapunov exponents of surrogated vowels.

母 音	音 声(bit)	生成音声(bit)
/a/	+0.001	+0.017
/i/	-0.048	+0.033
/u/	-0.062	-0.024
/e/	+0.007	+0.056
/o/	+0.091	+0.028

Table 5 The Lyapunov dimensions of surrogated vowels.

母 音	音 声(bit)	生成音声(bit)
/a/	1.01	1.03
/i/	0.00	1.15
/u/	0.00	0.00
/e/	1.01	1.09
/o/	1.16	1.06

次に、リアプノフ次元の計算結果をTable 3に示す。この結果より、全ての算出値が非整数次元となるので、これらが自己相似性をもっているといえる。以上のリアプノフ解析結果は、生成音声カオスであり、生成モデルにおけるカオスの特性が、声道の形状に大きく依存していることを示している。

より詳しく生成音声のカオスであることを検証するために、我々はサロゲーション法¹⁵⁻¹⁷⁾を用いた音声データの解析を行った。サロゲーション法は、ある時系列データが「カオスであるか、ノイズであるか？」という疑問を解決するために用いられる方法で、統計的解析における仮説検定の一種である。今回我々は、元の音声データから新たに時間的には相関をもつような（有色化された）データを生成し、再度この新たなデータに対しリアプノフ解析を行った。もし、元の音声データが有色化されたランダムデータであれば、新たに作り出したデータの性質は変化しない。逆に、新たなデータのリアプノフ解析結果が前の結果に比べて有意な差が認められれば、元の音声データが有色化されたランダムデータであるという仮説を棄却できる。すなわち、音声のゆらぎがノイズでないと検証される。

我々は元の音声データに対し、フーリエ変換を施し、パワースペクトラムの振幅情報を保持したまま、位相情報をランダム化した新たなデータを作成した。これらのデータに対する再度のリアプノフ解析結果をTable 4 およびTable 5に示す。これらのリアプノフ解析の結果は、最大リアプノフ指数は母音によって異なるが平均的に減少しており、リアプノフ次元は全体的に減衰し、サロゲーションによってカオスの特性が弱められることを示している。これらの変化は、有意な差であると考えられるので上記の仮説は棄却され、音声および生成音声はカオスであると推定される。

最後に、本モデルの生成音声および、石坂のモデルの生成音声をオーディオ機器によって再生した主観的な結果について述べる。石坂のモデルの生成音声は周期波形であるため、不自然さが感じられる。これに対して、本モデルであるカオス音声生成モデルの生成音声は、石坂のモデルより自然であるという結果を得ている。

7. 結 論

本研究において、我々は従来の音声生成モデルに対して、非線形な振動項を指数関数型で導入し、カオス特性をもつゆらぎを含んだ音声生成モデルを構成した。また、生成音声のゆらぎをもつことを、波形やアトラ

クタから視覚的に確認し、リアプノフ解析を用いてカオス特性の検証を行った。さらに、サロゲーション法の導入によって、生成音声のもつゆらぎが有色ノイズではないことを検証した。

本モデルで導入された非線形パラメータである ϵ_1 、 ϵ_2 の制御はさらに人間の声帯振動に近い非線形振動を提供できると考えられる。また、声道の形状が異なるとカオス特性が変化することから、声道形状と生成音声のカオス特性の関係を議論する必要がある。

音声生成モデルに関する研究分野では、サンプル音声データとその音声の発声時の声道形状を用いて、どの程度定量的に模倣した音声が可能かを検証する必要がある。

参 考 文 献

- 1) Lorenz, E.N., : Deterministic nonperiodic flow, J. Atoms. Sci. 20, p.130, (1963).
- 2) Mandelbrot, B.B., : *Fractal Geometry of Nature*, Freeman, San Francisco, (1982).
- 3) 山口達也, 永野倫法, 中川匡弘: 音声におけるフラクタル性に関する研究 (母音のフラクタル次元), IEICE, SP92-130, p.757, (1993).
- 4) Sano, M., Sawada, Y., : Measurement of the Lyapunov spectrum from a chaotic time series, Phys. Review. Lett., 55, 10, p.1082, (1985).
- 5) Eckmann, J.-P., Karmphorst, S. O., Ruelle, D., Ciliberto, S., : Lyapunov exponents from a time series, Phys. Rev. A, 34, 6, p.4971, (1986).
- 6) Wolf, A., Swift, J. B., Swinney, H. L., Vastano, J. A., : Determining Lyapunov exponents from a time series, Physica 16 D, p.285, (1985).
- 7) Sabanal, S., Nakagawa, M., : The Fractal Properties of Vocal Sounds and Their Application in the Speech Recognition Model, Chaos, Solitons & Fractals Vol.7, No.11, p.1825, (1996).
- 8) Flanagan, J.L., Landgraf, L., : Self-oscillating source for vocal-tract synthesizers, IEEE Trans. Audio and Electro-acoust., AU-16, p.57, (1968).
- 9) Flanagan, J.L., : *Speech Analysis Synthesis and Perception*, 2nd ed., Springer, New York, (1972).
- 10) Ishizaka, K., Flanagan, J.L., : Synthesis of Voiced sounds From a Two-Mass Model of the Vocal Cords, Bell System Tech.J., 51, 6, p.1233, (1972).
- 11) Hashiba, M., Ifukube T., : A Role of Waveform Fluctuation on Naturality of Vowels, Proc. J. Acous. Soc. Jpn, Fall Meeting, p.345, (1988).
- 12) 山口達也, 中川匡弘: 音声のフラクタル性とその評価法, IEICE, SP93-74, DSP93-75, p.79, (1993).
- 13) 古賀博之, 中川匡弘: カオス音声生成モデル, IEICE., NLP98-62, p.25, (1998).
- 14) Koga, H., Nakagawa, M., : A Chaotic Synthesis Model of Vowels, Proc. of The Fourth Int. Symp. on Artificial Life and Robotics, Oita, Japan, 19-22, p.556, (1999).

- 15) Theiler, J., : Some comments on the correlation dimension of $1/f^\alpha$ noise, *Physics Letters A*, 155, (8,9), p.480, (1991).
- 16) Prichard, D., Theiler, J., : Generating Surrogate Data for Time Series with Several Simultaneously Measured Variables, *Physical Review Letters*, 73, 7, p.951, (1994).
- 17) 池口徹：カオスと時系列解析について, 電気学会研究会資料, IP-95, 42, p.109, (1995).
- 18) 多谷虎男：不規則振動解析, 学会出版センター, (1981).
- 19) Fant, A., : *Acoustic theory of speech production's*—Copenhagen, Mouton & Co, (1960).
- 20) Takens, F., : Detecting Strange Attractors in Turbulence, in *Dynamical Systems and Turbulence*, Lecture Notes in Mathematics, Springer, 898, p.366, (1991).