

顔に見えるものとのインタラクションを実現するシステム 「かおさがし」の開発

松本 遥子[†] 堤 孝広[†] 寺澤 玲緒[†] 宮田 直貴[†] 藪 慎一郎[†] 宮田 一乗[‡]

[†]北陸先端科学技術大学院大学 知識科学研究科

[‡]北陸先端科学技術大学院大学 知識科学教育研究センター

A system Realizing Interaction with Face-looked Things, “KAOSAGASHI”

Yoko Matsumoto[†] Takahiro Tsutsumi[†] Reo Terazawa[†] Naoki Miyata[†]
Shinichiro Yabu[†] and Kazunori Miyata[‡]

[†]JAIST, School of Knowledge Science

[‡]JAIST, Center for Knowledge Science

{yk.mtmt, ttsutsumi, r.terazawa, s0850033, yabu.sh, miyata}@jaist.ac.jp

概要

身の回りには、顔に見えるものが多く存在する。もしそれらの顔が表情を変え、話をすることができれば、さらに楽しい発見や体験が行えるのではないだろうか。筆者らは、このような顔に見えるものとのインタラクションを目的とした作品「かおさがし」を提案し、新たに感情モデルと合成音声を加えたシステムを製作した。人形型デバイスで顔に見えるものを撮影し、人形の顔を触ったり、人形を揺らしたりすることによって、顔や音声インタラクティブに変化する。このシステムによって、顔に見えるものとのインタラクションが可能となった。今後、携帯電話向けアプリケーションや玩具として応用することによって、新たなエンタテインメントを提案することができる。

Abstract

There are a lot of face-like objects and materials. If those faces move, speak, and interact with us, it will be very amusing for many people, especially children. We therefore propose a system, which enables people to interact with those faces for entertainment. We have previously proposed a facial recognition system and a facial animation system to produce artwork which encourages people to interact with these face-like objects. We have newly updated the system, adding emotion modeling and synthetic voices. Users take a picture of a face-like thing with the doll-shaped device. Then, photograph take appearance in the face of the doll. When users touch the face of the doll or shake the body of the doll, the facial expression and the voice of the doll changes. We encourage people to interact with face-shaped objects for entertainment. We hope this system will be implemented in mobile gadgets such as mobile phones, and expect people, especially children, to use the gadget and play with faces in their environment.

1. はじめに

身の回りには、図1に示すように顔に見えるものが多く存在する。最近では、顔に見えるものを発見する驚きと、その顔のユニークさが注目を集め、顔に見える画像を集めたブログ[1]や写真集[2]も出版されている。また、デジカメやカメラ付き携帯の普及によって、発見したものをすぐに写真におさめられるようになったことも、身近な顔に目が向くようになった一因であると考えられる。

物が顔に見えるとそれがまるで生き物のように感じられる。しかし、その顔はあくまで物であり、その顔が表情をつくったり話をしたりすることはできない。もし、彼らが笑ったり泣いたり、話をすることができたりすれば、さらに楽しい発見や体験が行えるのではないだろうか。

本論文では、このような顔に見えるものとのインタラクションを目的とした作品[3]を提案する。つづいて、その作品の評価実験の結果を反映させ、新たに感情モデルと合成音声を加えた試みを実装したシステムについて述べる。



図1 身近な「顔」の例

2. 関連研究

人形やぬいぐるみ、おもちゃのロボットなど、生き物でないものを対象に会話をしたり遊んだりすることを子どもはよく行う。また、1997年にイグ・ノーベル賞を受賞するなど話題に上がった「たまごっち」など、電子ペットと呼ばれるカテゴリのおもちゃも多数市場に出回っており、人以外のモノを人に見立ててコミュニケーションを取ることは、ごく当り前の行為としてとらえられる。

一方、顔画像認識に関する研究は、歴史も古くこれまで数多く行われてきた[4][5][6]。最近では、笑顔を認識してシャッターが切れるデジタルカメラ[7]や、携帯電話での顔認識エンジン[8]、顔から年齢を推測し適切な広告を表示するシステム[9]、自分の顔と有名人との類似度を示すアプリケーション[10]、顔認証によるたばこの自動販売機などが開発されている。しかし、人の顔以外のものを顔として扱った研究は極めて少ない。藤原ら[11]は、車フロントマスクを顔に見立て、その線画に対するイメージを、表情と年齢印象において定量化・特徴比較した。しかし、身の回りのあらゆる物を顔として認識し、その顔に表情付けを行った研究は前例がなく、さらに、体験者の動作によって表情が変化するインタラクティブシステムも実現されていない。

我々は、そのようなシステムを実装し、顔に見える「物」と遊ぶという新たな体験を提案する。

3. システム概要

本作品では、小型カメラやタッチパネルディスプレイ、スピーカー、加速度センサを内蔵した人形型デバイスをインターフェースとして用いる。作品の使用例を図2に示す。

まず、体験者は身の回りから顔に見えるものを探す。図1に示すような偶然顔に見えるものを始めとして、物を顔のように並べたもの、顔を描いたイラスト、実際の人の顔など、目と口のように見えるものを有する被写体であればどのようなものでも良い。見つけた顔を、人形型デバイスの顔となるような位置に合わせ、背面のカメラで撮影する。撮影した画像はPCに転送され、目と口となる3領域を顔のパーツとして抽出する。顔を構成しているパーツの大きさや配置によって、あらかじめ設定した25種類のキャラクターの中から最適なキャラクターを決定し、そのキャラクターに依存した表情付けのアニメーションを人形型デバイスのディスプレイに表示する。同時に、キャラクター毎に異なる音声をスピーカーから出力する。これによって、体験者が見つけた顔が、人形の顔となってしゃべり始めたように見える。さらに、顔が表示されているタッチパネルディスプレイに触れたり、人形を揺らしたりすることによって、顔の表情や発言がインタラクティブに変化し、体験者を楽しませる。

撮影ボタンを再度押すと、顔とのインタラクションを終了する。同時に、撮影した顔が人型をしたキャラクターとなって動き回るアニメーションを背景画像に合成し、プロジェクタからスクリーンに投影する。これによって、顔が意思を持って動き回っているように見え、さらにキャラクターを体験者に強く印象付けることを目指した。



図2 作品の使用例

4. システム詳細

システム構成を図3に示す。システムは、1) 様々なパーツを内蔵した人形型デバイスと、2) 画像・音声処理用のPC、3) 今までに作られたキャラクターを表示するプロジェクタとスクリーンの3モジュールで構成される。

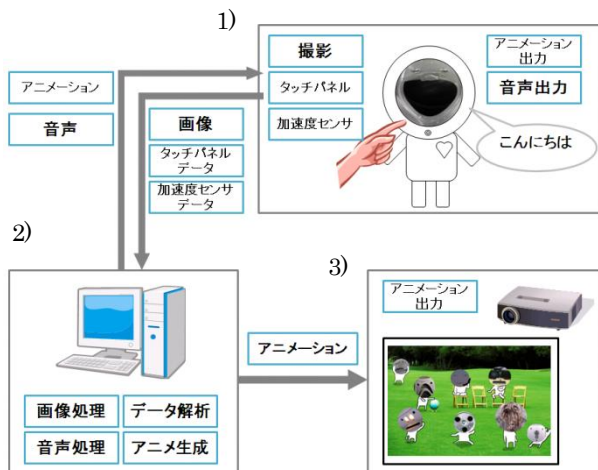


図3 システム構成

閾値と画素値を比較することで、背景が顔パーツの候補かを判断し、その結果を二値化した画像で表す。

最後に、稼働する領域を広げることでアニメーション時の動作を大きく見せるために、各領域を図7に示すように3ピクセル拡大し(3ピクセルの設定は経験則)、最終的な候補領域とする。これは特に面積の小さな顔パーツに対して有効な処理である。



図6 撮影された画像



図7 候補領域に分けられた画像

4.1 人形型デバイス

人形型デバイスは、背面に小型化カメラ、前面にタッチパネルディスプレイとスピーカ、胴体部分に撮影ボタンと加速度センサが人形型の合皮のカバーに内蔵されている。前面のディスプレイは、カメラに映る映像を出力し、体験者が撮影ボタンを押すと、PCからのアニメーション出力に切り替わる。人形型デバイスの展開図を図4に、カバーからデバイス内部を取り出した図を図5にそれぞれ示す。

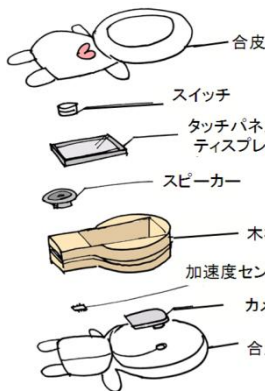


図4 人形型デバイス展開図



図5 デバイス内部

4.2.2 顔の各部位との対応

候補領域を抽出した後、それぞれの部位を顔の部位となる目、口、鼻へと対応させる。被写体は、マスキングされた領域が顔の輪郭として捉えられるように撮影するので、あらかじめ設定した座標値から顔の部位をある程度推測することができる。

図8に示すように、各部位との対応付けを各候補領域の重心座標によって決定する。上部に位置するものが目(図中の青と黄)に対応し、下部が鼻(図中の緑)、口(図中の赤)の候補領域となる。各領域は、面積の大きい順番にラベル付けし、ラベル順に判定して対応付けを行うことで、各部位の候補領域のうち、最大面積のものをひとつ選択する。

対応した部位の候補となる領域がない場合は、図9に示すように、あらかじめ設定した領域を部位の候補として対応付けする。この処理で、壁のような均一で部位候補の判定が困難な被写体に対しても部位を必ず対応付けることができる。



図8 候補領域と各部位との対応 図9 予め設定された領域(口)

4.2 顔の認識

撮影した画像をPCに転送し、顔の口と目となる領域を検出する。本節では撮影された画像から顔に見えるものを検出する方法について説明する。その後、感情モデルのパラメータ取得について述べる。

4.2.1 候補領域の取得

撮影した画像から、目、口、鼻の候補となりうる領域の判定を行う。画像処理にはOpenCVを応用した。カメラには図6に示すように円形のマスクを掛けて撮影し、画像中央の領域を対象に判定を行う。

画像はグレースケールに変換後、画像のノイズを少なくするために平滑化を行い、大津の手法[12]を用いて閾値を決定する。

被写体をアニメーションさせるために各部位に制御点を設ける。制御点は、図10に示すように、領域の左右両端の2点と横幅を4等分に分割する線分上の領域の上下端点である。これらの8個の制御点を各部位に同様に設定する。

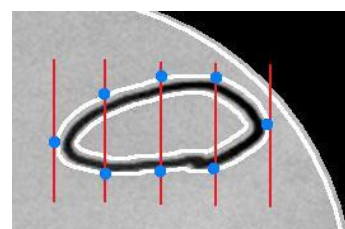


図10 アニメーションのための8個の制御点

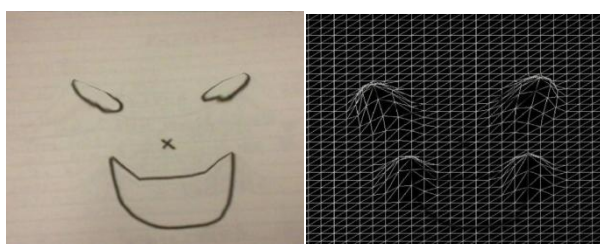
4.3 アニメーション

顔画像のアニメーション処理では、まず画像全体に対して 32×32 の分解能でメッシュを生成し、その上にテクスチャとなる顔画像を貼り付ける。そして認識処理で得られた目、鼻、口それぞれの端点の位置に最も近いメッシュの頂点を制御点として動かす事で、顔の表情をリアルタイムアニメーションとして表現する。また、各パーツのメッシュの制御点を動かす際に、移動による変形の影響が制御点の周りのメッシュにも及ぶようにする。すなわち、制御点からの距離に反比例してメッシュの頂点を移動する。図 11 に顔画像とそのワイヤーフレーム表示を示す。また、図 12 には変形後のものを示す。



(a) 顔画像 (b) ワイヤーフレーム表示

図 11 変形前の顔画像



(a) 顔画像 (b) ワイヤーフレーム表示

図 12 変形後の顔画像

各顔パーツの制御点の移動は以下のように行う。はじめに、制御点移動量の目標値(goal)とそれに達するまで要する時間(time)を設定する。そして、所要時間が 0 になるまで式(1)の一連の手順を計算し、算出される x をアニメーション中の出力値とする。

$$\begin{aligned} dx &= \max(dx_{old} - grad/e, \min(dx_{old} + grad/e, (goal - x) * e / (e + time))) \\ x &= x + dx \\ dx_{old} &= dx \\ grad &= (goal - x) / time \end{aligned} \quad \dots(1)$$

4.4 キャラクタ付けと音声処理

システム側で設定する 25 種類のキャラクタを表 1 に、キャラクタを決定するパラメータを表 2 に示す。ここで、表 2 のパラメータは、各部位に対して 1 つだけ 1 に、残りは 0 になるものとする。キャラクタは、目と口の大小、形によって決定する。表 2 に示したパラメータの各組み合わせ 80 通りに、表 1 に示したキャラクタをそれぞれ設定した。たとえば、目が小さくて丸く、口が大きくて上向きであれば、「あいどる」というキャラ

クタに決定する。表情の変化の例を図 13 に示す。

表 1 設定したキャラクタ

いなかっぺ	らっぱー	せいぎのみかた	だんてい
おとぼけ	やくざ	がいきくじん	うちゅうじん
おとこのこ	おかま	はずかしがりや	ふしぎちゃん
おんなのこ	びびり	くいしんぼう	かんさいじん
しゅうどう	ぎやる	あくのていおう	ねぼすけ
あいどる	いぬ	つんでれ	なきむし
えどっこ			

キャラクタによって、アニメーションや音声はそれぞれ異なる。音声の高さや速さ、発音、言い回しなどの異なる声をあらかじめ録音しておく。キャラクタが決定されると、そのキャラクタに対応する音声セットを選択し、発音に応じた表情の動きと連動して、人形型デバイス前面のスピーカから出力する。

表 2 キャラクタを決定する顔のパラメータ

顔の部位	要素	顔の部位	要素
目の大きさ	大	口の大きさ	大
	小		小
目の形	細い	口の形	水平
	丸い		斜め
	釣り目		口角が上向き
	垂れ目		口角が下向き
			丸い



図 13 キャラクタ付けと表情の変化

5. 評価と考察

2008年9月13-14日に日本科学未来館で開催されたIVRC2008東京大会にて本作品の展示を行い、その際に体験者66名にアンケートを実施した。対象者の男女比、年齢比を図14に示す。

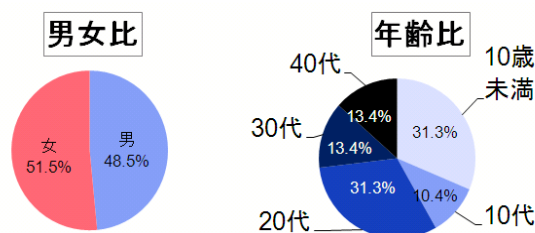


図 14 アンケート対象者の男女比と年齢

アンケートは、図15-17に示す質問に対して、5段階評価により回答してもらった。

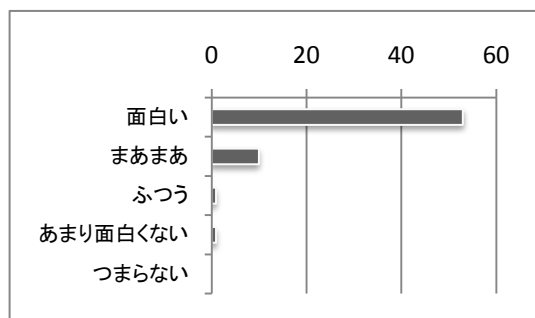


図 15「体験は面白かったか？」

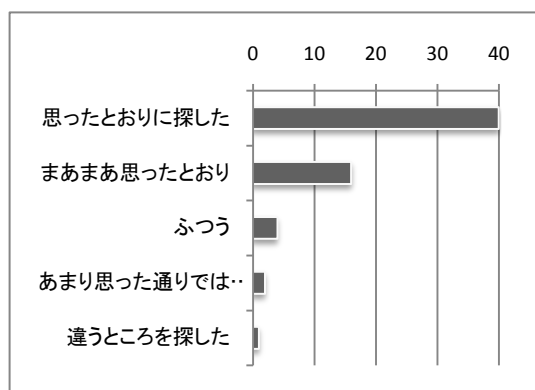


図 16「カメラは目や口を、思ったとおりに探してくれたか？」

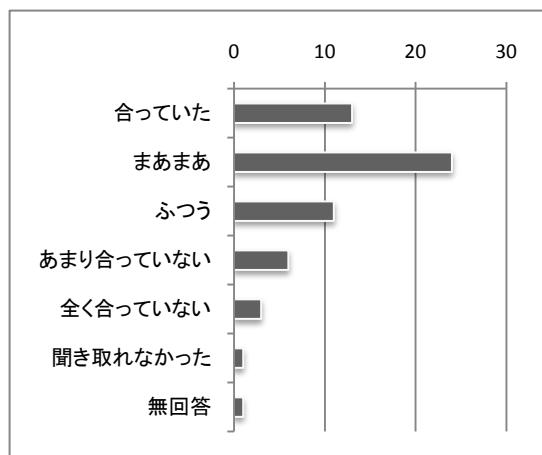


図17「声は、顔のイメージに合っていたか？」

また、自由記述により、以下のようなコメントを得た。

- ・ 発想が面白い、野外で試したら面白そう。
- ・ 友達と遊んでいるようだった。
- ・ 作った後、そのキャラともっと遊べばいい。
- ・ 顔を作った後のインタラクションが単調
- ・ いろんなキャラがあるので何回でも楽しめそう。
- ・ キャラクタの数が多。

図15が示すように、体験者の95.5%が、「面白かった」また

は「まあまあ面白かった」と評価した。この結果より、本作品によって顔に見えるものとのインタラクションが可能となり、新たな楽しみを提案することができたと考えられる。システムについては、図16より、カメラはほぼ体験者の思ったとおりに顔を探していると分かる。しかし、図17の「声は顔のイメージに合っていたか？」という質問については、「合っていた」または「まあまあ合っていた」と回答した体験者は56%に止まった。さらに、「いろんなキャラクタがあるので何回でも楽しめそう」というキャラクタに対する肯定的な意見がある一方、「キャラクタの数が多

い」というコメントも見られ、音声とキャラクタ付けに関しては、改良の余地があることが分かる。

また、インタラクションについても、単調であるという指摘があり、再考の必要があると考えた。

6. 感情モデルと音声合成による改良

5章での評価実験の結果を踏まえて、特に出力部である顔のキャラクタ設定法と音声処理のモジュールの改良を試みた。本章では、その改良法を述べる。

6.1 感情モデル

表情付けのために、4章では被写体を25種類のキャラクタに分類する手法を提案したが、撮影した顔があたかも生きているように見せるため、新たに図18に示す感情モデルを構築した。

このモデルは、表4に示す性格パラメータと表5に示す感情パラメータを内部に持つ。外部から表6に示す刺激が与えられ、そのパラメータにより表7に示すさまざまな反応を返したり、感情パラメータの値が変化して表情を変えたりする。

まず顔が撮影されたときに、抽出された顔の各部位から感情モデルのためのパラメータを決定する。使用するパラメータは表3に示す9個である。パラメータは各部位の重心座標、傾き、幅から決定する。但し、目の要素に対しては、右目と左目の平均値をそのパラメータとする。なお、各パラメータは0～1のレンジで正規化を行う。

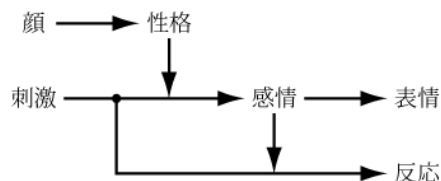


図 18 感情モデル

表 3 感情モデルの顔のパラメータ

顔の部位	要素
目	釣り目(角度)
	たれ目(角度)
	横幅の長さ
	両目間の距離
口	口角の上向き
	口角の下向き
	縦幅の長さ
	横幅の長さ
その他	目と口の距離

表4 性格のパラメータ 表5 感情のパラメータ

性格	感情
社交性	驚き
興奮しやすさ	恐怖
攻撃性	嫌悪
明るさ	怒り
動作の速さ	幸福
	悲しみ

表6 刺激一覧

刺激	要素
タッチパネル	額
	目
	鼻
	頬
	口
	顎
加速度センサ	一定以上の揺れ

表7 反応一覧

反応
笑う
目を閉じる
顔をしかめる
歯を食いしばる
くしゃみ
顔をぶるぶる震わす

撮影後は、性格と刺激から感情パラメータの変化量を感情モデルにより算出し、また感情と刺激から返す反応を決める処理を一定間隔で行う。刺激が与えられた場合はその刺激に対する反応によってアニメーションを行い、反応が特に無い場合は感情によって異なる表情のアニメーションを行う。

各パラメータは0から1までの値をとる。ただし刺激パラメータのみ0か1かの2値である。性格と刺激から求める感情変化は、中間の0.5よりも大きい場合に正の変化をし、小さい場合には負の変化をする。また、感情パラメータは0を初期値としており、感情変化が無い場合は0に収束するようにしている。

時刻 t での感情パラメータを $\bar{a}_t$ 、感情変化を $\Delta \bar{a}_t$ とすると、変化後の感情パラメータの値は式 (2) で計算される。

$$\bar{a}_{t+1} = (1 - 1/\tau)\bar{a}_t + k(\Delta \bar{a}_t - 0.5) \quad (2)$$

ただし τ, k は定数で正の値である。

各パラメータを求めるには3層パーセプトロンを用いている。その教師データとしてそれぞれ 50 以上のデータを用意し、バックプロパゲーションアルゴリズムを用いて学習を行った。

「怒り」「喜び」の感情パラメータ値が高い場合の表情の変化の例を図19に示す。

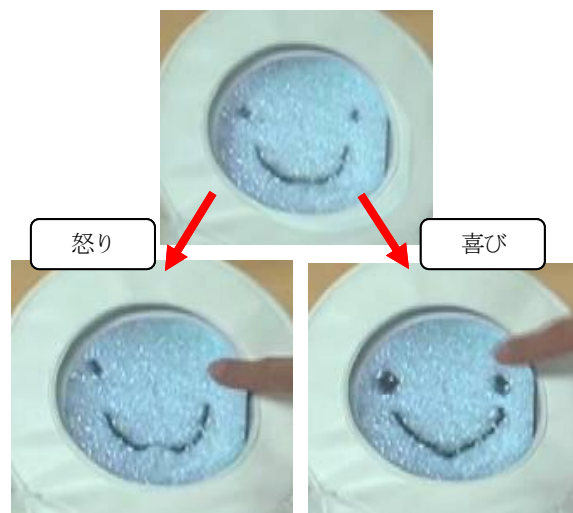


図19 撮影した顔(上)とインタラクションによる表情の変化(下)

6.2 音声処理

顔に見えるものに生き物らしさを持たせるもう一つの試みとして、声道モデルに基づいた音声をリアルタイムに合成する。刺激が入力された場合やしばらく入力がない場合に、リアルタイムで音声を合成して出力する。

6.2.1 音声生成モデル

図20に音声生成モデルを示す。音源となる三角波または雑音を、声道調音フィルタを通すことで音声を合成する[13]。このときパラメータとして、有声/無声の選択と、ピッチ、振幅、声道形状を指定できる。有声音は音源として三角波を、無声音は雑音を用いる。また、声道形状を指定することで母音を選択できる。

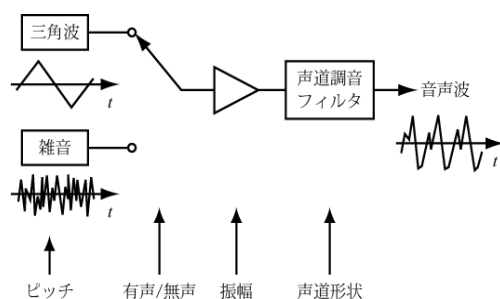


図20 音声生成モデル

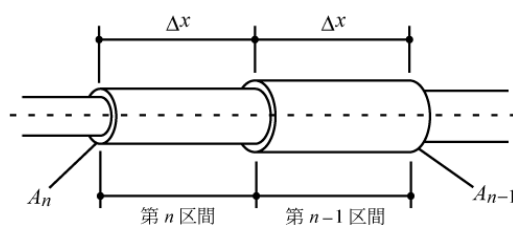


図21 微小音響管モデル

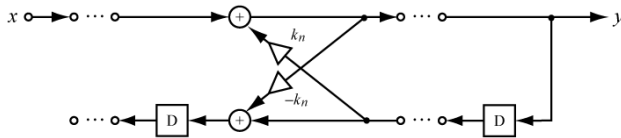


図 22 声道調音等価フィルタ

声道調音フィルタは声道を 1 次元音響管と考へ、そこから導かれた等価フィルタを用いる。この音響管モデルでは声道を断面面積が連続的に変化するものとし、分布定数系として扱う。

図 21 に示すように、声道を長さ Δx の区間に分け、各区間が一定の断面面積を持つものとして近似する。このモデルでは鼻腔を含まず、口腔のみで近似している。また、声道内のさまざまな要因による損失が、十分小さいものとして無視する。

この音響管のモデルから、振動媒体の体積速度と音圧に関する 1 次元の波動方程式を解いてフィルタとして表現すると、図 22 のようになる。フィルタの係数 k_n は、各区間の断面面積 A_n から式(3)で計算される。ただし $A_0=1$ とする。

$$k_n = (A_{n-1} - A_n)/(A_{n-1} + A_n) \quad (3)$$

6.2.2 韻律の制御と発話内容の生成

音声生成モデルを用いて、より話し声らしい音声を生成するために、イントネーションやアクセントをつける。改良後の作品では日本語の高低アクセントに基づく生成方法を用いている。

この方法は図 23 に示すように、各文節に割り当てられたアクセント型による高低と、文末へ行くほど低くなっていくピッチを足し合わせて高低のパターンを生成する。さらに、これに音素固有の声の高さが加わる。

発話内容は次の手順で生成する。まず、決められた範囲でのランダムな数の文節列を生成し、各文節にある範囲のランダムな数のモーラを生成する。そして各文節にアクセント型を割り当てる。このアクセント型には図 24 に示す日本語のアクセント型を用いる。

最後に、各モーラにランダムに母音や子音を割り当てる。

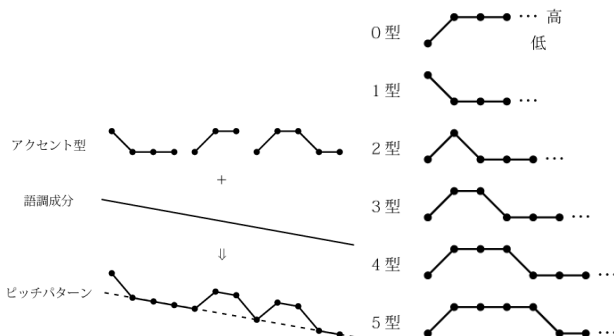


図 23 ピッチパターンの生成

図 24 日本語のアクセント型

6.2.3 感情モデルとの関係

以上の手順で音声を合成するときに、6.1 で示した感情モデルによってパラメータを変化させる。

設定できるパラメータは以下の通りで、各値は感情の値の関数として表現される。

- ・ 音量
- ・ ビッチ
- ・ 単位時間あたりの発話モーラ数
- ・ 抑揚の大きさ
- ・ 各文節間の間の長さ
- ・ 各文節の語尾の長さ
- ・ 最大文節数

音声は、タッチパネル、加速度センサからの入力を行うと、感情モデルによって音量・ビッチ・抑揚などを変化させ、表情の動きと連動して出力する。

6.3 改良により期待される効果

5 章で述べた結果では、改良前のシステムでは、顔とキャラクタ・声のイメージが合っておらず、また、キャラクタの数は多いが、インタラクションは単調であるという評価を得た。

そこで、改良後のシステムでは、人間を模した多くのキャラクタに顔を当てはめるのではなく、「顔たち」だけが話す言葉としてリアルタイム合成音声を用い、身の回りにある顔自体の個性を引き出すことを目指した。また、感情のみを伝えることによって、体験者が意味を推測できる余地を残した。さらに、感情モデルによって、体験者の操作が人形の感情を変化させ、顔との対話性を高めることを目指した。

7. 改良後のシステムの考察

7.1 改良後の展示と評価

2008 年 11 月 7-8 日に各務原テクノプラザで開催された IVRC2008 岐阜大会、同年 11 月 22-24 日に石川県産業展示館で開催されたいしかわ夢未来博にて、改良後の作品の展示を行った。体験者を観察すると、改良前のシステムとの明らかな違いとして以下のような様子が見られた。

- 体験者が人形の音声に対し、「何を話しているのか」と耳を近づける様子が減った。
- 体験者の操作（タッチパネルに触れる／人形を揺らす）の回数が増えた。

A)より、体験者がリアルタイム音声合成に対して違和感を持たずに操作を行っていることが分かる。B)は、話している内容に意味を持たせないことで、人形が話し終えるのを待たずに次の操作を行う体験者が増えたこと、操作回数によって表情が次々変化するようシステムを変更したことが原因と考えられ、改良前のシステムに比べ、体験者が積極的にインタラクションを行うようになったことが分かる。

改良したシステムの詳細な評価に関しては、今後、アンケートをとるなど、改良前のシステムとの比較を行う予定である。

7.2 考察

改良前、改良後の作品とも、10 歳以下の体験者が長時間に渡って、または繰り返し来訪して遊ぶことが多かった。また、顔に見えるものだけでなく、体験者が描いた顔のイラストや、体験者自身の顔を撮影し、インタラクションを行うことによって、周囲の体験者と笑いあい、楽しむ様子が見られた。

一部の体験者からは、「玩具にして欲しい」、「コミュニケーションツールとして使えるのではないか」という感想も得られ、展示のための作品ではなく、携帯電話向けアプリケーションや玩具のような持ち歩けるツールとして発展させることにより、さらに利用の可能性が広がり、楽しい体験を提供できると考える。

8. おわりに

本論文では、顔に見えるものとのインタラクションを実現するシステム「かおさがし」について述べた。本システムによって、顔に見えるものとのインタラクションが可能となり、また、新たな楽しみを提案することができたと考えられる。これは、すべてのものに魂が宿っているというアニミズムの考え方を具体化したものとも考えられる。

今後は、さらに楽しめるインタラクションを追及するとともに、携帯電話向けアプリケーションなどとしての開発を行い、どこでも顔を探し、遊ぶことができるシステムを実現したい。

謝辞

研究開発助成をいただいたIVRC実行委員会に深謝いたします。本研究の一部は、文部科学省「大学院教育改革支援プログラム」の助成をいただきました。

参考文献

- [1] <http://kaodarake.blog86.fc2.com/>
- [2] 龍山 悠一, “Who are you?”, 光村推古書院, 2002.
- [3] 松本 遥子, 堤 孝広, 寺澤 玲緒, 宮田 直貴, 藪 慎一郎, 宮田 一乗: かおさがし: 顔に見えるものとのインタラクション, エンタテインメントコンピューティング2008講演論文集, pp.149-150, 2008.
- [4] Rainer Lienhart and Jochen Maydt: “An Extended Set of Haar-like Features for Rapid Object Detection”, Proceedings of the 2002 IEEE International Conference on Image Processing, Vol.1, pp.900-903, 2002
- [5] 石井洋平, 本郷仁志, 丹羽義典, 山本和彦: “顔検出のための特徴量とその領域の検出”, 日本顔学会誌, Vol.3, No.1, pp.33-41, 2003
- [6] 平山 高嗣, 岩井 儀雄, 谷内田 正彦, “顔画像認識を用いた施錠セキュリティシステムFACELOCKの開発”, 電気学会論文誌C, Vol. 124, No. 3, pp.784-797, 2004
- [7] SONY 製サイバーショットのスマイルシャッター機能など <http://www.sony.jp/products/Consumer/DSC/DSC-T200/feat1.html>
- [8] オムロン: OKAO VISION, <http://www.omron.co.jp/ecb/products/mobile/>
- [9] NEC: デジタルサイネージソリューション http://www.nec.co.jp/solution/video/d_signage.html
- [10] ジェイマジック: かおちえき! <http://www.j-magic.co.jp/j/kaocheki/>
- [11] 藤原 孝幸, 興水 大和, 川澄: “車フロントマスクの顔表情と年齢印象の研究”, CGAC 2007予稿集, S3-2, pp.1-6, 2007.
- [12] 大津展之, “判別および最小2乗基準に基づく自動しきい値選定法”, 電子通信学会論文誌, Vol.J63-D, No.4, pp.349-356, 1980.
- [13] 古井貞熙, “音声情報処理”, 森北出版, 1998.

松本 遥子



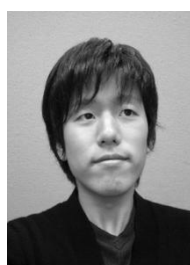
2008 年金沢工業大学情報フロンティア学部メディア情報科卒業。同年より、北陸先端科学技術大学院大学知識科学科博士前期課程在学。

堤 孝広



2008年東京工業大学・情報工学科卒業。同年より、北陸先端科学技術大学院大学知識科学科博士前期課程在学。

寺澤 玲緒



2008年東海大学開発工学部感性デザイン学科卒業。同年より、北陸先端科学技術大学院大学知識科学科博士前期課程在学。

宮田 直貴



2008年東京理科大学基礎工学部電子応用工学科卒業。2008年より北陸先端科学技術大学院大学知識科学科博士前期課程在学。

藪 慎一郎



2008 年石川工業高等専門学校卒業。同年より、北陸先端科学技術大学院大学知識科学科博士前期課程在学。

宮田 一乗



1986年東京工業大学大学院・総合理工学研究科・物理情報工学専攻修士課程修了。同年、日本アイビーエム(株)東京基礎研究所入社。1998年東京工芸大学芸術学部助教授。2002年より、北陸先端科学技術大学院大学知識科学教育研究センター教授。博士(工学)。コンピュータグラフィックスおよびデジタル映像表現に関する研究に従事。情報処理学会、芸術科学会、映像情報メディア学会、ACM、IEEE等会員。