

深層学習を用いた入力音声に適した顔表情生成

西村亮佑^{1,a)} 酒田信親² 富永登夢¹ 土方嘉徳³ 原田研介¹ 清川清²

概要: 近年, VR で利用する高品質な三次元 CG キャラクタに対する需要が高まっている. その中で, 表情を提示する顔のアニメーションは, コンテンツのリアリティを高めるうえで重要である. 現在, 表情アニメーションの生成手法は, カメラを用いた手法が一般的であるが, カメラのキャプチャ範囲に制限があることや, 顔の向きによっては人間の表情を取得できないなど, カメラの設置環境に依存する問題がある. そこで, 音声情報のみによる表情アニメーションを生成するシステムが有用であると考えた. 本研究は, LSTM-RNN を用いた先行研究を参考に, 音声から顔表情を生成し, それをコミュニケーションに利用可能なシステムにすることを目標とする. 音声の解析に人間の聴覚感度を考慮した A 特性手法を用いることで, 先行研究よりも表情の精度が向上することを定量的に確認した.

1. はじめに

近年, VR で利用する高品質な 3 次元 CG キャラクタに対する需要が高まっている. その中で, 表情を提示する顔のアニメーションは, コンテンツのリアリティを高めるうえで重要な要素を占めている. 現在, 表情アニメーションの生成手法は, カメラを用いた手法が一般的であるが [1][2], カメラのキャプチャ範囲に制限があることや, 顔の向きによっては人間の表情を取得できないなど, カメラの設置・撮影環境に依存する問題がある.

カメラを用いる手法に対し, 音声情報のみを用いて表情を推定する手法は, これらの問題を解決する可能性がある. 音声のみを用いた手法には, 音素の解析により口の形を再現する手法 [3][4] や, 深層学習を用いて顔全体の表情を再現する手法 [5][6] がある. しかし, これらの手法は生成できるアニメーションが顔の一部分に限定されていることや, 話の内容をスピーチのみしか扱っていないことから, 人間同士の対面コミュニケーション時に表される豊かな表情を推定可能であるかの検証は行われていない.

そこで本研究では, 音声のみを入力としてもっともらしい豊かな表情を生成するシステムの構築を目指す. 自然な表情を生成するためには, 口の動きだけでなく, 人間同士の対面コミュニケーションに使えるような, 頬の起伏や眉の動きにも焦点を当てる必要がある. そこで, リカレントニューラルネットワーク (RNN) を用いて表情を推測する

手法 [6] をベースとして, 豊かな表情を生成できるシステムを構築し, 提案システムの評価を行う.

2. 関連研究

本節では, センサや画像データ, 音声データから表情を推測し, 3D モデルに適用する研究を紹介する.

まず, 画像データから表情を推定する研究を紹介する. 3D モデルに表情を適用する際は, ブレンドシェイプに基づく方法が主に用いられる. [7][8][9]. ブレンドシェイプとは, 顔の一部分が無表情から変化して 3D モデルを複数合成することで, あらゆる表情を提示する手法である. あらかじめ用意する顔を決定する際には, 顔の動きを包括的に測定する Facial Action Coding System (顔面動作符号化システム) [10] を応用する. 2002 年に発表された改訂版 FACS には, 41 の顔の基本動作が定義されており, これらは顔の解剖学的な知見をもとに構成されている.

ブレンドシェイプのパラメータを決定するためには, 顔特徴データを取得する必要がある. 様々なセンサを用いた手法が存在する. 深度センサを用いて顔形状を取得している研究 [11][12] では, Kinect[13] を深度センサとして使用しているが, 表情の精度がデバイスに依存することや, 画角の制限が問題である. Zell らは, モーションキャプチャマーカを用いて, マーカと 3D モデルを対応付けることで表情を推定している [1][14]. この手法は精度が高い一方, マーカの取り付けや, 3D モデルとの位置合わせ, カメラの調節など, 多くの段階で手間を要する. 画像データのみから顔の特徴を推定する手法も存在する [15][16]. この手法は, 3D 顔形状データベースと, 顔の特徴点及び輪郭の

¹ 大阪大学大学院 基礎工学研究科 システム創成専攻

² 奈良先端科学技術大学院大学 先端科学技術研究科 情報科学領域

³ 関西学院大学 商学部

^{a)} nishimura@hlab.sys.es.osaka-u.ac.jp

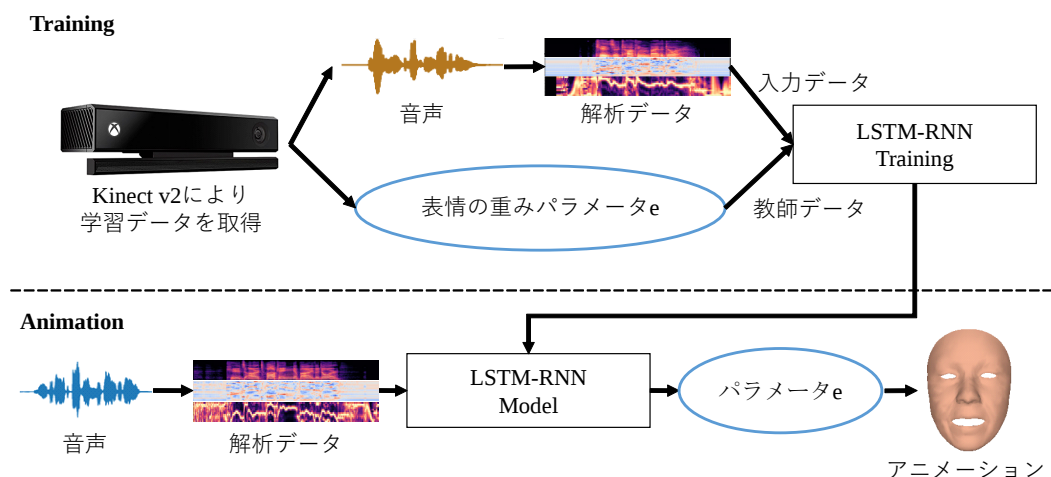


図 1 提案システムの全体像

2D 特徴点の真値を持つ画像データベース [17][18][19] を用いて、ランダムフォレストによる回帰分析を行うことで実現している [20][21]。顔の輪郭などに誤差が生じる場合があるが、画像のみで実現できる点が優れている。

ブレンドシェイプのパラメータの値は顔形状の違いにより変化するため、適切なパラメータを求めるためには任意の人物の 3D 顔形状も求める必要がある。事前の Kinect などの深度センサや Cyberware[22] などのレーザースキャンにより、複数人の 3D 顔形状のデータベースを構築する。任意の人物の 3D 顔形状を再構築する方法として、主成分分析 (PCA) に基づく線形モデルを用いる手法 [23][24] と、テンソルを用いた多重線形顔モデルを用いる手法 [17][25] がある。これらの手法を用いて、任意の人物の顔画像から特徴点を抽出することで、特徴点に最も一致した 3D 顔形状を再構築できる [2][20][21]。複数人の 3D 顔形状のデータベースに FaceWarehouse[17] がある。これは、150 人の様々な民族の、47 種類の FACS に基づく表情を備えるデータベースである。

続いて、音声データから表情を推定する研究を紹介する。音声データのみを用いた表情アニメーションの生成には、音響特性を用いた手法 [3][4][26] と深層学習を用いた手法 [5][6][27] がある。音響特性を用いた手法は、音の発生源である口周りの表情の動きを推定するには十分であるが、目や眉の動きなど、動きが音に影響しない部分の表情は推定できない。それに対し、深層学習を用いた手法では、入力に周波数解析したデータを用いており、音声と表情を対応付けることができる。しかし、音声解析を学習の入力とする際、音声解析データが多くなるほど計算コストがかかり、リアルタイム性に欠けるため、コミュニケーションに用いることが難しくなる。リアルタイムで用いるためには、どの音声解析データが表情生成のために有効か検証すべきである。また、入力データがスピーチや演技に限られているため、相槌や笑う・驚くといった反応が含まれる、

実際のコミュニケーションを想定していない。

音声データから表情を推定する研究を紹介したが、出力された人間同士の対面表情の評価は適切に行う必要がある。音声と表情は必ずしも一致しないことや、人の感情を音と表情のどちらから読み取るかは文化によって異なることが知られており [28]、これは真値と深層学習による出力が必ずしも一致しない可能性を示している。従って、出力として得られた表情が入力の音声に適した度合を表す新たな評価手法を提案する必要がある。

3. 提案システム

本稿では、音声データのみを入力とし、LSTM-RNN を用いて音声に適した表情を出力するシステムを提案する。学習の入力は音声解析データ、出力はブレンドシェイプの表情重みパラメータとする。人間同士の対面コミュニケーションに利用するシステムの全体像を図 1 に示す。

3.1 音声データ

機械学習に用いる音声データの解析方法について説明する。Pham ら [6] は音声波形データをスペクトログラム、ケプストラム、クロマグラムの 3 種類に解析し、入力データとした。本研究では、音声に A 特性手法による知覚的重みづけを行い、入力データとする。実装には、Python のライブラリである LibROSA[29] を用いる。

3.1.1 スペクトログラム

スペクトログラムとは、音を周波数解析したものである。形式は様々であるが、横軸を時間、縦軸を周波数として表すことが多い。音声のスペクトログラムでは、特に母音に特徴が出やすいため、音声の解析には有効な表現方法である。出力の例を図 2 に示す。

3.1.2 ケプストラム

ケプストラムとは、信号のスペクトログラムをフーリエ変換したもので、スペクトログラムの細かな変動となら

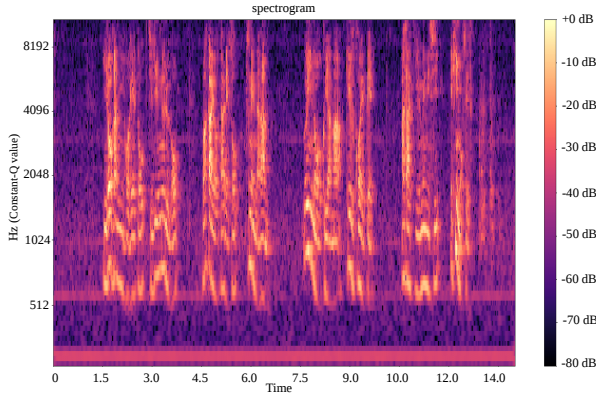


図 2 スペクトログラム

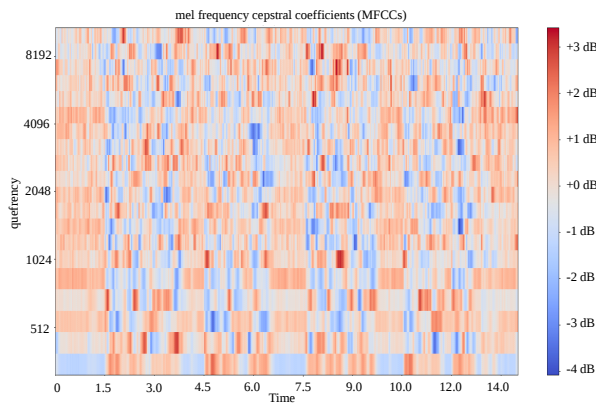


図 3 ケプストラム

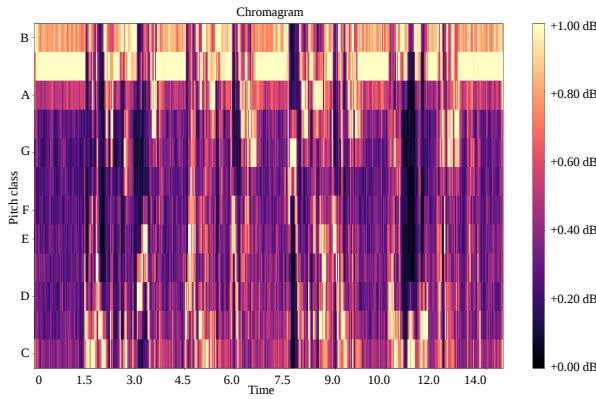


図 4 クロマグラム

かな変動を分離することができる。スペクトログラムは母音に特徴が出やすいため、その変動は音声解析するうえで重要な情報であると考えられる。出力例を図 3 に示す。

3.1.3 クロマグラム

クロマグラムは全周波数帯域のエネルギーを、[C, Db, D, Eb, E, F, Gb, G, Ab, A, Bb, B] の 12 音階に落とし込んだものである。音声にメロディの特徴がある場合は、クロマグラムのデータが特徴を持つことになる。出力の例を図 4 に示す。

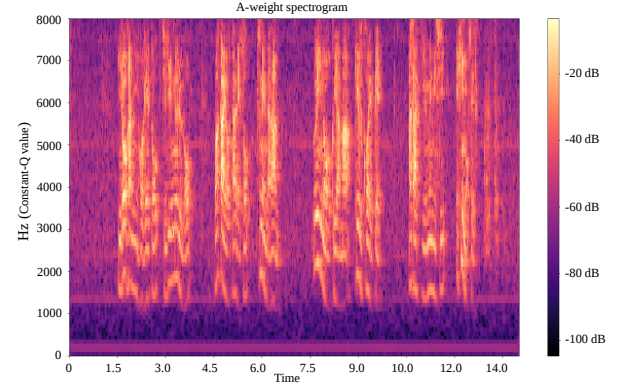


図 5 A 特性

3.1.4 A 特性

A 特性とは人間の聴覚の特性を考慮して、周波数に知覚的重みづけを行ったものである。人間の音声は、高音域や低音域にはあまり特徴が含まれないため、この手法を用いることが有効だと考える。出力の例を図 5 に示す。

3.2 顔特徴データ

本稿では、深度センサとして Kinect v2 を用いて顔特徴データを取得し、ブレンドシェイプの手法を用いて表現する方法について述べる。

3.2.1 Kinect v2 による顔特徴の取得

Kinect for Windows SDK に含まれる HDFace[30] を用いることで、3 次元で 1346 個の特徴点を取得することが可能である。Pham ら [6] の研究で用いられているデータセットは感情に焦点を当てているが、コミュニケーションで現れる表情については考慮されていない。そのため、Kinect v2 を用いて、目的を達成するために必要なデータセットを作成する。

図 6 に HDFace で得られる顔の特徴点を示す。HDFace は、事前に骨格を識別することで、顔の位置をロバストに認識出来る。Kinect v2 で得られる顔の特徴点データはノイズが大きいため、ノイズ除去フィルタを施す必要がある。フィルタには [12] に用いられている指数移動平均を用いる。 i フレーム目の特徴点の行列を $\mathbf{t}_i$ と置くと、フィルタをかけた後の行列 $\mathbf{t}_i^*$ は

$$\mathbf{t}_i^* = \frac{\sum_{j=0}^k w_j \mathbf{t}_{i-j}}{\sum_{j=0}^k w_j} \quad (1)$$

$$w_j = e^{-j \cdot H \cdot \max_{l \in [1, k]} \|\mathbf{t}_i - \mathbf{t}_{i-l}\|} \quad (2)$$

で表される。ここで、 k はフィルタに用いるフレーム数、 H は実験的に決定する定数である。本研究の実験では、 $k = 5$ 、 $H = 0.01$ とした。このフィルタにより、大きなノイズの除去に成功したが、依然として揺らぎのようなノイズが残る。このノイズは、ブレンドシェイプの手法で近似することで取り除くことができる。

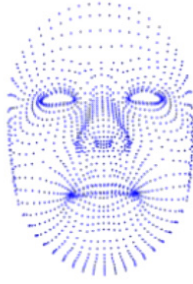


図 6 HDFace で得られる顔の特徴点



図 7 本研究で用いたブレンドシェイプ

3.2.2 ブレンドシェイプによる表情生成

Kinect により得られた顔形状を、ブレンドシェイプを用いて表現する。この手法により、顔の頂点数が異なる 3D モデルでも、同様の顔形状をあらかじめ用意することで、同じ表情を再現できる。

今回用意しておく表情は、図 7 に示す 22 個の顔形状とした。口の動きは正確に再現すべきであるため、16 個用意した。また、左右で顔の形状に差がある方が人間らしさが表現できると考えたため、左右別の動きをした顔形状を用意した。

ブレンドシェイプを用いた顔形状は以下の式で表される。

$$S = B_0 + \sum_{i=1}^m (B_i - B_0) \cdot e_i \quad (3)$$

ここで、 B_i はあらかじめ用意しておいた顔形状、 m は B_i の個数、 $e_i \in [0, 1]$ はブレンドシェイプの重みパラメータである。Kinect v2 の顔形状に一致するように、ブレンドシェイプの重みパラメータ e_i を算出し、学習データの出力に用いる。

3.3 機械学習の構成

音声の時系列に伴った信号であることから、LSTM-RNN による機械学習を行うことで、音声から表情を推定する。

入力データには 3.1 節で述べた音声解析データを用いる。教師データには 3.2 節で述べた表情の重み e とする。

今回用いる LSTM-RNN のフレームワークでは、音声データ x_t を入力とし、表情の重み e_t のみを出力とする。ここで、 $t = 1, \dots, T$ であり、 T はフレーム数である。任意の時間 t において、学習モデルは入力特徴ベクトル x_t から e_t を推定する。表情の重み e は $[0, 1]$ 内に制限されているため、ReLU の活性化を用いて出力が非負となるようにする。学習モデルは二乗誤差が最小になるように訓練する。

$$E = \sum_t \|y_t - \hat{y}_t\|^2 \quad (4)$$

3.4 実装

ここでは、学習データを取得するための実装について述べる。学習の入力に用いる音声データは、LibROSA[29] を用いて解析し、学習の出力に用いるブレンドシェイプの重みパラメータは、Kinect から取得したデータから算出し、表情の真値として用いる。

ブレンドシェイプのパラメータ算出は、線形の最適化問題となる。ブレンドシェイプを用いた表情は、式 (3) で表現でき、 e_i が求めたいパラメータとなる。最適化問題は以下のように表される。

$$\arg \min_{e_j} E[e_j] = \arg \min_{e_j} \|V - S\|_2 \quad (5)$$

この最適化問題の勾配ベクトルは、

$$\frac{\partial E[e_j]}{\partial e_j} = -E[e_j]^{-1} \cdot \sum_{i=1}^n (V - S)_i \cdot (B_j - B_0)_i^T \quad (6)$$

で与えられる。ここで、 n は顔の頂点の個数である。今回は HDFace を用いているため、 $n = 1346$ となる。これらの式を用いて、 $e_i \in [0, 1]$ の条件を満たすようにパラメータ e_i を算出する。最適化アルゴリズムで、今回の要件を満たす手法は、L-BFGS-B と TNC と SLSQP がある。すべて試用した結果、今回は最も実行時間が短く、精度の良かった TNC を採用する。

得られた最適化の結果を図 8 に示す。誤差が現れる部分は、瞬きの際のまぶたと、下唇周辺が多い。コミュニケーションを取ることを考慮して、これらの誤差は許容範囲内とみなす。

4. 実験

3 節で、音声データのみを入力とし、LSTM-RNN を用いて音声に適した表情を出力するシステムを提案した。本節では、提案システムの評価を行う。

データセットには、遠隔地でのコミュニケーションなど、本システムの利用シーンに即したデータとするため、電話の会話データを用いた。顔形状と音声データを同時に取得することで、音声と表情にずれが生じないようにして

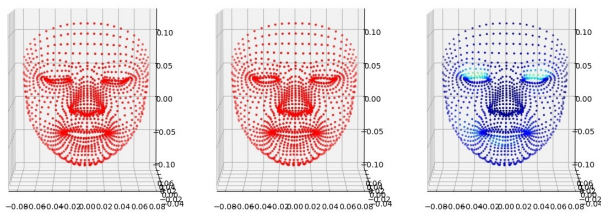


図 8 ブレンドシェイプのパラメータ算出の結果. 左: Kinect により得られた顔形状. 中: 最適化により得られたパラメータを用いてブレンドシェイプにより構築した顔形状. 右: 左と中の顔形状の誤差

いる. Kinect v2 のマイクから取得できる音声は, サンプル周波数が 16kHz, ビットレートが 32bit IEEE float である. サンプル周波数が低めだが, 人間の声の特徴は 8kHz 以下に多く現れるため, 利用可能である.

実際には, 男性 2 名が電話で会話するシーンの合計 186,142 フレーム (約 104 分) のデータセットを用いる. 学習の入力データとして, 以下の 3 パターンを用いる.

- (1) 音声の元データ
- (2) ノイズ除去を施したデータ
- (3) (1)+(2) にホワイトノイズを追加したデータ

ノイズ除去を行う理由は, ノイズを特徴として学習させないためである. (3) で, (1) と (2) とノイズ加算のデータを用いている理由は, 音声を加工することにより, データセット数を増やすことができるためである.

学習の構造は, 中間層を 2 層とし, 1 層目は 600 次元, 2 層目は 600 次元とする. 深層学習のライブラリは Microsoft の CNTK を用いる. 学習環境は, CPU: Core-i5-6500 3.20GHz, GPU: GTX 1080, RAM: 16GB である.

5. 結果と考察

4 節で挙げた 3 パターンの入力データに対する出力の評価を行う.

(1) や (2) では, 依然として顔形状が震える印象が残る. この対策として, 現在は音声情報を A 特性手法のみを用いているが, 先行研究に用いられたデータを使って再度検証する余地がある. また, データセットが不足している可能性も考えられる.

(3) では, 無音の時に震える問題を改善できている. これは, ノイズ加算による影響や, 音声データを修正したため, データセットが増えたことによる影響だと考えられる. 音声に対する推定表情は, まだ正確とは言えないため, データセットをさらに増やすことが必要だと考えられる.

今回出力された表情をコミュニケーションで利用可能かどうかの検証はまだできていない. そこで, 被験者実験での検証を予定している. 被験者実験では, 実際にシステム

を体験した際に人がどう感じるかの評価を行う. カメラを用いて人間の表情をキャプチャし, 3D モデルの表情を動かせるシステムを用いた際に, 任意の期間を音声で推定した表情に切り替えると, 使用する人は気付くかどうかの評価を提案する.

システムの評価基準に以下の 2 つを挙げる.

- 初見で違和感が生じた期間の正解率
 - 複数回見て, 最終的に違和感が生じる期間の正解率
- 一つ目の評価はシステムの評価である. リアルタイムでこのシステムを使用する場合, 同じ表情は 1 度しか見ることができないため, そこで違和感が生じるかどうかの評価ポイントとなる. 二つ目の評価は表情の推定精度の評価である. じっくり見て違和感が生じなければ, 表情の推定精度は高いとみなすことができる. システムの評価式を以下に示す.

$$\text{推定表情期間の正解率} = \frac{\text{推定表情期間} \cap \text{違和感}}{\text{表情推定期間}} \quad (7)$$

$$\text{表情のあいまいさ} = \frac{\text{推定表情期間以外} \cap \text{違和感}}{\text{表情推定期間以外}} \quad (8)$$

表情に対する違和感は個人差があると想定されるため, 被験者の人数を多く設定する必要がある. この評価では, カメラを用いて 3D モデルの表情を動かしている期間に, 違和感が生じる可能性がある. これは, 人間の表情に対するイメージのあいまいさの評価に繋がると予想できる. また, このシステムを用いることで, 人間は様々な表情を取ることができるが, コミュニケーションを取る際には, どれほどの表情が必要かが示唆出来ると想定している.

6. おわりに

本稿では, 音声の入力により, 表情アニメーションを生成するシステムの提案を行った. RNN を用いて表情を推測する手法 [6] を参考に, 音声に A 特性手法による知覚的重みづけを施すことで, 表情推定精度を向上した. また, Kinect を用いてコミュニケーションに対応したデータセットを作成する方法を述べ, その学習結果についての考察を述べた. 今後の研究方針は, 現在取り組んでいるシステムの修正, および被験者実験によりシステムの評価を行うことである.

参考文献

- [1] Zell, E., Lewis, J., Noh, J., Botsch, M. et al.: Facial retargeting with automatic range of motion alignment, *ACM Transactions on Graphics (TOG)*, Vol. 36, No. 4, p. 154 (2017).
- [2] Cao, C., Hou, Q. and Zhou, K.: Displaced dynamic expression regression for real-time facial tracking and animation, *ACM Transactions on graphics (TOG)*, Vol. 33, No. 4, p. 43 (2014).
- [3] Edwards, P., Landreth, C., Fiume, E. and Singh, K.: JALI: an animator-centric viseme model for expressive

- lip synchronization, *ACM Transactions on Graphics (TOG)*, Vol. 35, No. 4, p. 127 (2016).
- [4] Taylor, S., Kim, T., Yue, Y., Mahler, M., Krahe, J., Rodriguez, A. G., Hodgins, J. and Matthews, I.: A deep learning approach for generalized speech animation, *ACM Transactions on Graphics (TOG)*, Vol. 36, No. 4, p. 93 (2017).
 - [5] Karras, T., Aila, T., Laine, S., Herva, A. and Lehtinen, J.: Audio-driven facial animation by joint end-to-end learning of pose and emotion, *ACM Transactions on Graphics (TOG)*, Vol. 36, No. 4, p. 94 (2017).
 - [6] Pham, H. X., Cheung, S. and Pavlovic, V.: Speech-driven 3d facial animation with implicit emotional awareness: a deep learning approach, *The 1st DALCOM workshop, CVPR* (2017).
 - [7] Pighin, F., Szeliski, R. and Salesin, D. H.: Resynthesizing facial animation through 3d model-based tracking, *Computer Vision, 1999. The Proceedings of the Seventh IEEE International Conference on*, Vol. 1, IEEE, pp. 143–150 (1999).
 - [8] Lewis, J. and Anjyo, K.-i.: Direct manipulation blendshapes, *IEEE Computer Graphics and Applications*, Vol. 30, No. 4, pp. 42–50 (2010).
 - [9] Li, H., Weise, T. and Pauly, M.: Example-based facial rigging, *Acm transactions on graphics (tog)*, Vol. 29, No. 4, ACM, p. 32 (2010).
 - [10] Ekman, P. and Friesen, W.: Facial action coding system: a technique for the measurement of facial movement, *Palo Alto: Consulting Psychologists* (1978).
 - [11] Zhang, Z.: Microsoft kinect sensor and its effect, *IEEE multimedia*, Vol. 19, No. 2, pp. 4–10 (2012).
 - [12] Weise, T., Bouaziz, S., Li, H. and Pauly, M.: Realtime performance-based facial animation, *ACM transactions on graphics (TOG)*, Vol. 30, No. 4, ACM, p. 77 (2011).
 - [13] Xbox: Xbox One 用 Kinect, <https://www.xbox.com/ja-JP/xbox-one/accessories/kinect>.
 - [14] Seol, Y., Lewis, J., Seo, J., Choi, B., Anjyo, K. and Noh, J.: Spacetime expression cloning for blendshapes, *ACM Transactions on Graphics (TOG)*, Vol. 31, No. 2, p. 14 (2012).
 - [15] Cao, X., Wei, Y., Wen, F. and Sun, J.: Face alignment by explicit shape regression, *International Journal of Computer Vision*, Vol. 107, No. 2, pp. 177–190 (2014).
 - [16] Ren, S., Cao, X., Wei, Y. and Sun, J.: Face alignment at 3000 fps via regressing local binary features, *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pp. 1685–1692 (2014).
 - [17] Cao, C., Weng, Y., Zhou, S., Tong, Y. and Zhou, K.: Facewarehouse: A 3d facial expression database for visual computing, *IEEE Transactions on Visualization and Computer Graphics*, Vol. 20, No. 3, pp. 413–425 (2014).
 - [18] Huang, G. B., Ramesh, M., Berg, T. and Learned-Miller, E.: Labeled faces in the wild: A database for studying face recognition in unconstrained environments, *Technical Report 07-49, University of Massachusetts, Amherst* (2007).
 - [19] F.Tarrés, A.: GTAV Face Database, <http://gps-tsc.upc.es/GTAV/ResearchAreas/UPCFaceDatabase/GTAVFaceDatabase.htm>.
 - [20] Cao, C., Weng, Y., Lin, S. and Zhou, K.: 3D shape regression for real-time facial animation, *ACM Transactions on Graphics (TOG)*, Vol. 32, No. 4, p. 41 (2013).
 - [21] Pham, H. X., Pavlovic, V., Cai, J. and Cham, T.-j.: Robust real-time performance-driven 3D face tracking, *Pattern Recognition (ICPR), 2016 23rd International Conference on*, pp. 1851–1856 (2016).
 - [22] Cyberware: Rapid 3D Scanners, <http://cyberware.com/>.
 - [23] Blanz, V. and Vetter, T.: A morphable model for the synthesis of 3D faces, *Proceedings of the 26th annual conference on Computer graphics and interactive techniques*, ACM Press/Addison-Wesley Publishing Co., pp. 187–194 (1999).
 - [24] Paysan, P., Knothe, R., Amberg, B., Romdhani, S. and Vetter, T.: A 3D face model for pose and illumination invariant face recognition, *Advanced video and signal based surveillance, 2009. AVSS'09. Sixth IEEE International Conference on*, Ieee, pp. 296–301 (2009).
 - [25] Vlastic, D., Brand, M., Pfister, H. and Popović, J.: Face transfer with multilinear models, *ACM transactions on graphics (TOG)*, Vol. 24, No. 3, pp. 426–433 (2005).
 - [26] Cao, Y., Tien, W. C., Faloutsos, P. and Pighin, F.: Expressive speech-driven facial animation, *ACM Transactions on Graphics (TOG)*, Vol. 24, No. 4, pp. 1283–1302 (2005).
 - [27] Zhou, Y., Xu, S., Landreth, C., Kalogerakis, E., Maji, S. and Singh, K.: VisemeNet: Audio-Driven Animator-Centric Speech Animation, *arXiv preprint arXiv:1805.09488* (2018).
 - [28] 高木幸子, 平松沙織, 田中章浩: 日本人の顔と声による感情表現の収録とその評価 (マルチモーダル, 感性情報処理, 視知覚とその応用, 一般), 電子情報通信学会技術研究報告. HIP, ヒューマン情報処理, Vol. 111, No. 283, pp. 51–56 (2011).
 - [29] GitHub: librosa/librosa Python library for audio and music analysis, <https://github.com/librosa/librosa>.
 - [30] Pterneas, V.: HOW TO USE KINECT HD FACE, <https://pterneas.com/2015/06/06/kinect-hd-face/>.