

物体運動からの任意照明・姿勢画像の生成

中島 朗子 牧 淳人

(株) 東芝 研究開発センター

E-mail: {akiko.nakashima , atsuto.maki}@toshiba.co.jp

概要

一定の照明条件下で姿勢変化する物体を固定カメラで撮影した少数の画像から、任意の照明下で同じ対象物体が任意の姿勢をした画像を生成する手法を提案する。この手法では、姿勢を固定した物体の任意照明下での画像を線形和で表す画像基底の概念を用いるが、この基底生成のために物体を固定することや、照明条件を変化させる設備は必要としない。代わりに、物体の姿勢が異なる入力画像間で対応点探索を行い、得られる対応付けに従って画素を並び替えることにより、画像基底を生成する。特に、対応点探索に線形結合係数の繰り返し計算を導入し、初期の対応付けが不十分な場合にも安定して画像基底が求まることを示す。また、対応点探索に伴って得られる三次元形状の利用により、任意の姿勢に対しても画像基底が求まることも示す。実験により提案手法の有効性を検証する。

Synthesizing Pose and Lighting Variation from Object Motion

Akiko NAKASHIMA and Atsuto MAKI

TOSHIBA CORPORATION, Corporate Research & Development Center

E-mail: {akiko.nakashima , atsuto.maki}@toshiba.co.jp

abstract

We present a novel method to synthesize images of a 3D object in arbitrary poses illuminated from arbitrary directions, given a few images of the object in unknown motion under static lighting. Our scheme is underpinned by the notion of *illumination image basis* which spans an image space of arbitrary lighting, and we propose to generate it by a recursive search for the correspondence between input images and by subsequent realignment of their pixels. Using the 3D surface of the object that also becomes available in this procedure, we synthesize images of the object in arbitrary poses while arbitrarily varying the direction of lighting by combination of the illumination basis images. The effectiveness of the entire algorithm is shown through experiments.

1 はじめに

三次元物体の見え方は照明条件や物体の姿勢によって変化する．これは三次元物体の画像を扱う技術において避けられない問題であり，既にこの問題に対処するための様々な提案がなされている．例えば，三次元物体認識では，物体の姿勢が変化したり照明位置が変化した場合でも安定して物体を認識するための手法が数多く提案されている [8, 1, 14]．しかしながら，これらの手法では，様々な物体姿勢や様々な照明条件下における画像をあらかじめ登録しておく必要がある．そのためには，大量の画像を撮影しなければならないが，そのような環境変化を登録時に全て再現することは現実的ではない．少ないフレーム数の画像から様々な条件下における物体の見え方を生成することが望まれる．これは，物体認識だけではなく，例えば文献 [11] で提案されているようなレンダリングの前処理等にも利用することができる．

任意照明下での画像は，異なる方向から照明を当てた数フレームの画像の重ね合わせで表すことができる [10, 3]．このような重ね合わせに使われる画像は基底画像とよばれ，その生成手法は，入力画像を撮影する際の照明と物体の位置関係によって以下の2つに分類することができる．

- (a) 姿勢を固定した物体に対して照明位置を変化させる場合
- (b) 固定された照明下において物体が姿勢変化する場合

(a),(b) 共に，物体と照明の相対的な位置関係は同じであることから，物体表面上の輝度変化としては同等の効果をもつことがわかる．(a) の環境で撮影された入力画像を用いる場合 [4, 10]，物体の姿勢は変化しないため，画像間の対応付けを計算する必要はないが，照明位置を変化させるために，照明リグ等の大掛かりな設備が必要となる．また，人物やペット等を撮影対象とする場合，被験者は撮影中静止していなければならない，ストレスを受けることになる．このようなストレスやコストを軽減するために (b) の条件で撮影された入力画像を用いる手法が提案された [9, 13]．この手法では，撮影対象が静止



図 1: 異なる方向から照明を当てて撮影された入力画像 (上段) と照明変動画像基底 (下段)．

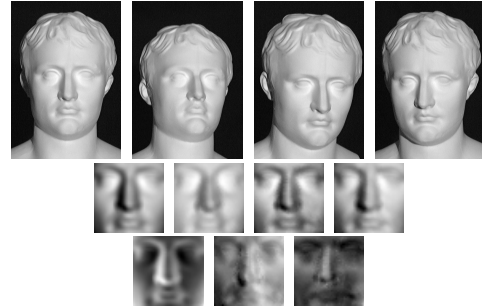


図 2: 固定照明下で姿勢変化する物体の入力画像 (上段)，画素並び替え画像 (中段)，照明変動画像基底 (下段)．

することも大掛かりな照明設備も必要とせず，固定カメラを1台用意するだけでよい．その代償として，対象物体が姿勢変化するために必要な画像間の対応付けを計算によって求める．本稿では，撮影環境の簡略化と被写体の負担軽減を優先する立場から，(b) の条件で撮影された入力画像を扱うことにする．

文献 [9] の手法では，幾何輝度拘束 [5] に基づいて密な対応点探索を行い，得られた対応付けに従って画素を並べ替えることにより，あたかも姿勢を固定した対象物体に異なる方向から照明を当てて撮影したかのような画像を作り，基底画像を生成している．そのため，画像間の対応付けが安定して得られることが重要となる．また，姿勢変化については，入力画像におけるのと同じ姿勢に限られ，任意ではなかった．本稿では，以下2つの目的でこの手法を改良する．1つは，繰り返し計算を導入し，対応付けを安定して行えるようにすることである．もう1つの目的は，任意姿勢に対して照明変動画像基底を生成することである．実は，幾何輝度拘束に基づく画像間の対応付けは，対象物体表面の三次元形状を介して得られる．この形状情報を利

用して任意姿勢に拡張する.

まず, 第2節で照明変動画像基底の概念を述べる. 第3節で, 画像間の対応付けを安定に行うため, 照明変動画像基底の生成法を改良し, 実際に安定化することを実験的に示す. 第4節で, 任意姿勢に拡張できることを示す. 最後に第5節で本稿のまとめと今後の課題を述べる.

2 照明変動画像基底とは

本節では, 照明変動画像基底の概念について説明する. 無限遠にある点光源を考える. このとき, 任意方向から照射された対称物体の画像は, 異なる方向から照射した画像の重ね合わせによって表現できる [10, 3]. 例えば, 対象物体表面が完全拡散反射特性をもつ場合は, 3 フレームの線形結合で表現される [10]. 輝度値は必ず正の値であることを考慮すると, 自分自身の影が自分に落ちる self-shadow がない場合, 第 j フレームの画像 $\mathbf{I}(j)$ は,

$$\mathbf{I}(j) = \max([\check{\mathbf{I}}(1) \ \check{\mathbf{I}}(2) \ \check{\mathbf{I}}(3)]\mathbf{a}(j), 0), \quad (1)$$

で表される [2]. ここで, $\check{\mathbf{I}}(j)(j = 1, 2, 3)$ はそれぞれ基底画像を表し, ここではこれらをまとめて照明変動画像基底とよぶ. $\mathbf{a}(j)$ は線形結合係数を成分にもつ3次元ベクトルである.

照明変動画像基底は, 姿勢を固定した対象物体に異なる3方向から照明を照射して撮影した3フレームの画像として得ることができる. しかし, 実際には雑音が存在するため, 3フレームより多くの画像を撮影し, 主成分分析 (PCA) を適用し, その主成分ベクトル (固有ベクトル) として照明変動画像基底を得る¹. 照明変動画像基底の例を図1に示す. 上段の4フレームは, 様々な方向から照明を照らして撮影した石膏像の正面画像である. これらの入力画像に対してPCAを適用した結果を図1の下段に左から主成分 (固有値) の大きい順番に並べた. ここでは, PCAの効果が得られるための最低フレーム数である4フレームを用いている. 一方, 本稿の目的は, 固定した点光源下で姿勢変化する対象物体を固定カメラで撮影して得られる入力画像か

¹PCAの代わりに, 文献 [4] で提案されている手法を用いてもよい. 影を考慮した手法であるため, より精密な画像基底を得ることができる.

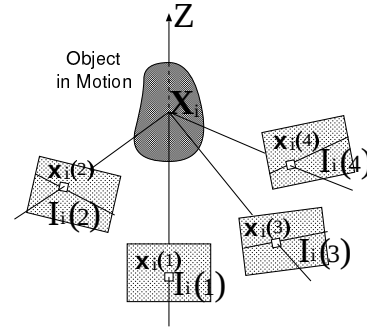


図 3: 幾何輝度拘束による画像間の対応付け. 基準フレームのある点に対して, 他フレームの対応点はエピポーラ線上に存在する. 奥行き Z を適当に与え, 輝度拘束を満たすような対応付けを探索する.

ら照明変動画像基底を生成することである. その一例を図2に示す. 最上段が入力画像, 後に説明する中段の処理を経て, 最下段が得られる照明変動画像基底である. 線形結合係数 $\mathbf{a}(j)$ を適当に決めれば, 基底画像の線形結合によって任意の照明下での画像を生成することができる.

3 照明変動画像基底の生成

姿勢変化する物体画像から照明変動画像基底を生成するためには, 画像中の画素を並び替えることにより, あたかも同じ姿勢で異なる方向から照明を当てたかのような画像を作る. これは, 図1の入力画像を作ることに相当する. このとき重要なことは, 画像間の画素の対応付けを正確に行うことである. 本節ではまず, 幾何輝度拘束 (Geotensity constraint) [5] について簡単に述べ, この拘束に基づいた照明変動画像基底の生成法 [9] について説明する. この生成法に繰り返し計算を導入し, 初期情報が少ない場合でも安定して対応付けを行い, 精密な基底を生成できるように改良する.

3.1 幾何輝度拘束 (Geotensity constraint)

カメラと光源の位置は固定され, その下で対象物体が姿勢変化するものとする. 第1フレームを基準画像とみなす. ここでは説明を簡略化するため, 弱中心射影モデル [7] を仮定する.

弱中心射影モデルは、カメラと対象物体との距離に比べて対象物体表面の奥行きが十分小さいものとした場合に成り立つ。

n_j フレームの画像を $\mathbf{I}(j) (j = 1, \dots, n_j)$ とする。ワールド座標系における i 番目の物体表面上の点 $\mathbf{X}_i = (X_i, Y_i, Z_i)^\top$ が第 j フレームの画像座標 $\mathbf{x}_i(j) = (x_i(j), y_i(j))^\top$ に射影される。第 1 フレームを基準とし正準化すれば、弱中心射影モデルにおいて

$$\mathbf{x}_i(j) = \mathbf{M}(j) \begin{pmatrix} \mathbf{x}_i(1) \\ Z_i \end{pmatrix} + \mathbf{t}(j) \quad (2)$$

の関係が成り立つ。ここで、 $\mathbf{M}(j)$ はスケールリングされた回転行列 $\mathbf{R}(j)$ の第 1, 2 行からなる 2×3 行列、 $\mathbf{t}(j)$ は平行移動を表す 2 次元ベクトルである。これらを合わせて運動パラメータとよぶことにする。式 (2) は画像 $\mathbf{I}(1), \mathbf{I}(j)$ 間のエピポラ拘束を表す。これは、基準画像における点 $\mathbf{x}_i(1)$ の対応点 $\mathbf{x}_i(j)$ が、運動パラメータ $\mathbf{M}(j), \mathbf{t}(j)$ によって定まる式 (2) のエピポラ線上に存在することを意味する。対応点がエピポラ線上に存在する様子を図 3 に示す。奥行き Z_i に依存して対応点はエピポラ線上を移動することがわかる。

ここで、画像上の点 $\mathbf{x}_i(j)$ の輝度 $I_i(j)$ を考える。ある Z_i を選べば、式 (2) によって画像間の対応付けが与えられ、画像 $\mathbf{I}(j)$ から輝度値 $I_i(j)$ を直接得ることができる。一方、式 (1) に従って他フレームの輝度から $I_i(j)$ を推定することもできる。以下では、こうして得られる輝度を区別するため、前者を観測値とよび $I_i(j)$ で、後者を推定値とよび $\hat{I}_i(j)$ で表すことにする。線形結合係数 $\mathbf{a}(j)$ と非零である輝度の観測値 $I_i(j)$ が得られれば、式 (1) における $\check{I}_i(j) (j = 1, 2, 3)$ は

$$\begin{bmatrix} \check{I}_i(1) & \check{I}_i(2) & \check{I}_i(3) \end{bmatrix} = \mathbf{I}_{n'_j}^\top \mathbf{a}_{n'_j}^\top (\mathbf{a}_{n'_j} \mathbf{a}_{n'_j}^\top)^{-1} \quad (3)$$

で与えられる。ここで、 n'_j は輝度の観測値 $I_i(j)$ が非零であるフレーム数、 $\mathbf{I}_{n'_j}$ は $I_i(j)$ を要素にもつ n'_j 次元ベクトル、 $\mathbf{a}_{n'_j}$ は $\mathbf{a}(j)$ を列に並べた $3 \times n'_j$ 行列である。この時、推定値 $\hat{I}_i(j)$ は

$$\hat{I}_i(j) = \max \left(\begin{bmatrix} \check{I}_i(1) & \check{I}_i(2) & \check{I}_i(3) \end{bmatrix} \mathbf{a}(j), 0 \right) \quad (4)$$

となる。

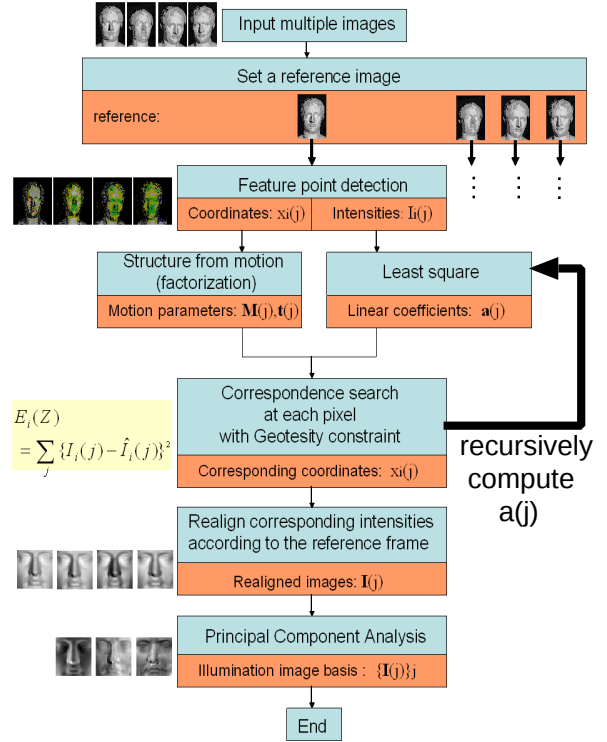


図 4: 照明変動画像基底の生成アルゴリズム。

推定値の誤差を、観測値との差の二乗和で定義する：

$$E_i = \sum_{j=1}^{n_j} (I_i(j) - \hat{I}_i(j))^2 \quad (5)$$

対応付けが正しければ、推定値 $\hat{I}_i(j)$ は観測値 $I_i(j)$ と一致する。即ち、 $E_i = 0$ が成り立つ。これが、幾何輝度拘束 (geotensity constraint) である。この拘束条件に基づいて、画像間の対応付けを行うことができる。あらかじめ 4 点以上の特徴点を検出しておけば、運動パラメータは因子分解法 [12] を特徴点座標に適用することによって得られ、線形結合係数は最小二乗法 (あるいは、特異値分解; SVD) を特徴点の輝度値に適用することによって得られる。こうして得られる画像間のパラメータは全画素に対して共通である。これらのパラメータを用いて、ある奥行き Z_i を適当に与えれば、対応点の輝度の観測値と推定値が得られ、誤差 E_i を計算することができる。実際には、観測値 $I_i(j) (j = 1, \dots, n_j)$ は雑音を含むため、誤差を最小にする奥行き Z_i をエピポラ線に沿って探索することになる。

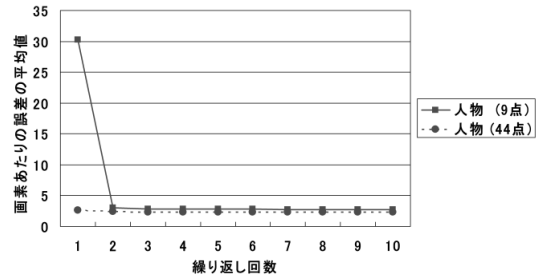
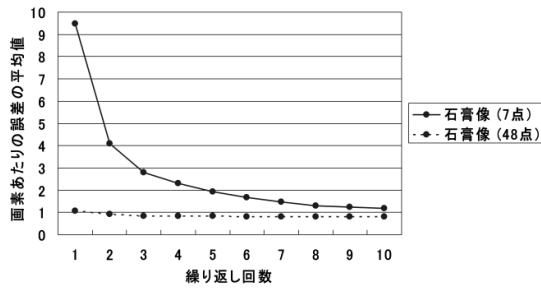


図 6: 繰り返し計算による 1 画素当りの輝度誤差の変化。

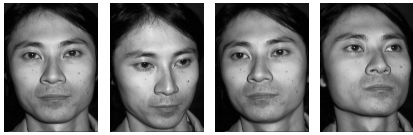


図 5: 入力画像. 固定照明下で姿勢変化する人物を撮影した。

3.2 基本アルゴリズム

照明変動画像基底を生成するアルゴリズムを図 4 に示す. 図 4 には本稿で提案する改良点も同時に示されているが, ここではまず, 細い矢印で表される基本アルゴリズム [9] について述べる.

まず, 入力画像の中から基準フレームを決定する. 入力画像間で対応する特徴点を 4 点以上検出する. 検出された疎な特徴点に対して, その座標に因子分解法を適用し運動パラメータを得, その輝度値にランク 3 の拘束をつけた SVD を適用し線形結合係数を得る.

これらのパラメータを使って, 幾何輝度拘束に基づき, 画像間の対応点探索を各画素毎に密に行う. 得られた対応付けに従って, 入力画像の画素を基準画像と同じ姿勢になるように並び替える. その例を図 2 の中段に示す. 最後に, 得られた並び替え画像に主成分分析 (PCA) を適用し, 固有値の大きい固有ベクトルを照明変動基底画像とする. 得られる基底の例を図 2 の下段に示す. 基準画像に別の入力画像を選べば, 同じアルゴリズムで基底を生成することができ, 入力画像に現れた姿勢変化に対して任意の光源環境を想定した画像を生成することができる.

3.3 対応点探索の安定化

先に述べたとおり, 基本アルゴリズムでは, あらかじめ対応付けされた疎な特徴点から運動パラメータと線形結合係数を計算し, それらを用いて密な対応点探索を行っていた. 従って, 照明変動画像基底の精度は, 特徴点の性質に依存する. ここで, 特徴点の対応付けは正規化相関などテンプレートマッチングによって求めるのが通常である. この時, 周囲の濃淡パターンがテンプレートに最もよく似ている点に対応する特徴点として検出される. 即ち, 対応点周囲の輝度パターンが大きく変化しないことが前提になっている. しかしながら, 本稿で扱う入力画像では対象の姿勢変動に伴い照明との相対位置がフレーム毎に異なるため, 対応点周辺での輝度パターンも異なっている. このような画像に対するテンプレートマッチングからは, 輝度パターンの変動が比較的小さな点が検出される. そうして検出された個々の特徴点が輝度の変動を十分表現していないと, 特徴点の数が少ない場合は特に, 線形結合係数を正しく推定することは困難となる. このような問題を解決するために, 本稿では, 図 4 の太い矢印で示される繰り返し計算を導入する. 一旦解いた密な対応点探索の結果を利用し, 改めて線形結合係数を解き直し, 対応点探索を繰り返し行うことによって, 対応付けの安定化を試みる.

対応点探索後に線形結合係数 $\mathbf{a}(j)$ を再計算する場合, 対応点探索によって得られる全対応点の中から, はずれ値を除去する必要がある. 初めに検出された特徴点が不十分である場合, 即ち, $\mathbf{a}(j)$ の初期値が正しくない場合, 間違って対応付けされる点が出てくるためである. 各点の信頼度は式 (5) で表される輝度の誤差を基準

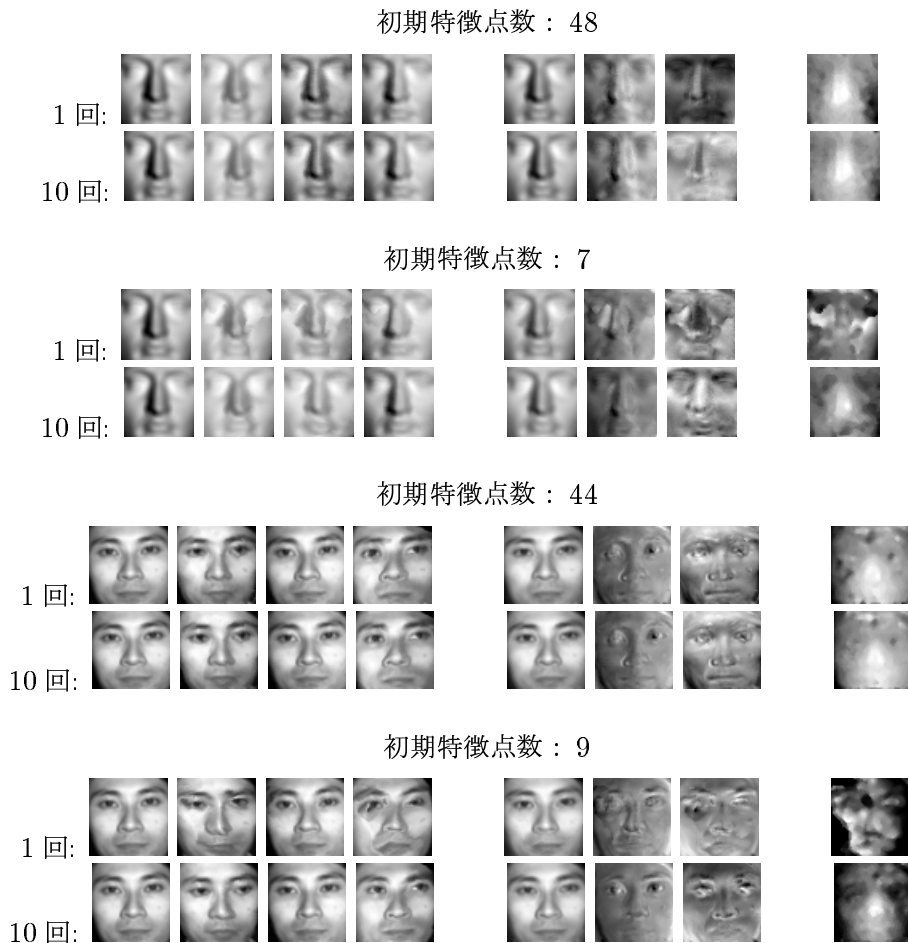


図 7: 繰り返し回数 1 回 (上段) と 10 回 (下段) の結果. 並び替え画像 (左) 照明変動画像基底 (中) 奥行きマップ (右).

に測ることができる. ここでは, 再計算する度に誤差の閾値を適応的に決め, その閾値より誤差が大きい点ははずれ値として再計算には用いないことにする. 残りの対応点の輝度値を使って, 次回の対応点探索に用いる線形結合係数 $\alpha(j)$ を計算する.

3.4 実験

以下では, 先に提案した繰り返し計算によって, 対応点探索を安定に行えることを実験的に示す. 対象物体として石膏像と人物顔を用いる. 入力画像を図 2 の最上段と図 5 に示す. 初期条件の違いによる安定性を調べるために, 初期の特徴点数を, 石膏像では 48 点と 7 点, 人物顔では 44 点と 9 点とした場合に対して, 提案手法を適用した.

まず, 式 (5) で定義される誤差の変化を図 6 に

示す. 横軸は繰り返し回数を, 縦軸は全画素に対する誤差の平均を表す. 実線は特徴点の少ない場合の変化, 点線は特徴点数の多い場合の変化を表す. 石膏像, 人物顔, いずれに対しても, 初期特徴点数の多い場合は, 誤差は徐々に減少し, 初期特徴点数の少ない場合は, 2 回目の誤差が急激に減少している. 前者は, 大量の初期特徴点から 1 回目の計算によって線形結合係数をほぼ正しく求めることができるためであると考えられる. 後者は, 特徴点数が少ないために 1 回目では線形結合係数を正しく求められないが, 2 回目には線形結合係数を計算するために用いた点の数が急激に増加し, 線形結合係数を比較的正確に求められるためであると考えられる. どの場合にも, 誤差は 10 回目にはほぼ収束していることがわかる.

1 回目の対応点探索の結果と 10 回目の対応点

探索の結果を図 7 に示す．左側から 4 フレームは得られる対応付けに従って正面向きに画素を並び替えた画像，次の 3 フレームは照明変動画像基底を示す．また，3.1 節で述べたように，対応付けは奥行き Z_i を介して行なわれる．対応付けの結果得られる奥行きマップを図 7 の一番右側に示す．（但し，奥行きは 256 階調で表されている．）初期特徴点数が多い場合は，基底，奥行き共に精度が大きく変化することはないが，奥行きの精度がわずかながら向上している．例えば，石膏像の左頬周辺や右目付近，人物顔の右頬周辺である．初期特徴点数が少ない場合は，1 回目の結果では 3 番目の基底画像や奥行きのくずれが顕著であるが，10 回目には基底，奥行き共に精度が向上している．尚，人物顔のいずれの場合も，目や額の辺りに陥没が見られるが，これは目が実際には非剛体であることや，額等ではてかり等の鏡面反射がおこりやすいためであると考えられる．

以上の結果は，繰り返し計算により精度が向上すること，特に初期特徴点の数が少ない場合に効果的であることを例証するものである．顔などでテクスチャの少ないものを対象とする場合，特徴点を大量に検出することは難しいため，このことは実用上有用であると考えられる．

4 任意姿勢画像の生成

3.2 節で述べたとおり，照明変動画像基底は，基準画像に現れる姿勢に対して生成することができる．基準画像を変更すれば同じ手順で異なる姿勢の照明変動画像基底を生成できた．逆に，物体姿勢の種類は入力画像と同じものに限られていた．本節ではこれを，入力画像のフレーム数を増やすことなく，任意の姿勢に拡張できることを示す．

4.1 三次元形状の利用

幾何輝度拘束に基づく対応点探索の結果，基底画像だけではなく三次元形状も同時に求まる．式 (1) に従って基底画像を足しあわせてできる任意光源下での画像を，得られる形状に対してレンダリングし，自由な回転を加え画像平面に射影するだけで，任意の姿勢における任意光源



図 8: 様々な顔向き画像の生成.

下での画像を生成することができる．光源方向の調節は線形結合係数で行う．

4.2 実験

図 7 の最下段に示した照明変動画像基底と奥行きを用いて，実際に任意の姿勢に対して任意の照明変動下における画像を生成できることを示す．奥行きマップで表される三次元形状を様々な角度で回転させ，照明変動画像基底の中の第 1 基底画像をテクスチャとしてレンダリングした結果を図 8 に示す．テクスチャは，式 (1) で $\mathbf{a}(j) = (1, 0, 0)^T$ としたことと同等である．形状を利用することによって，入力画像のフレーム数を増やすことなく，任意の姿勢に対する画像を実際に生成できることがわかる．

これらの中から入力画像とは異なる姿勢を一つ選び，あたかも照明位置を変化させたかのような画像を生成した．その結果を図 9 に示す．これらは，図 7 の最下段に示した 3 フレームの基底画像を式 (1) に従って足しあわせ，三次元形状にレンダリングし，回転を与えたものである．線形結合係数は任意に選ぶことができ，ここでは表 1 に示す値に従って変化させた．実際に，照明位置の変化を合成している様子が確認できる．

このような方法で任意姿勢の画像を生成するためには，形状の精度が大きく影響する．繰り返し計算によって形状の精度を高めることは，任意姿勢の画像を生成する上でも有用であることがわかる．



図 9: 入力画像には含まれない姿勢における様々な照明下での画像. 図 8 の最下段に示される照明変動画像基底と奥行きマップを用いて生成した.

表 1: 図 9 の画像を生成する際に用いた線形結合係数. 基底を求める際に得られる特異値を考慮した.

j	1	2	3	4	5	6	特異値
$a_1(j)$	1	1	1	1	1	1	$\times 13441$
$a_2(j)$	25	25	13	1	1	13	$\times 1765$
$a_3(j)$	25	1	1	1	13	25	$\times 996$

5 おわりに

本稿では, 固定照明下で姿勢変化する三次元物体を固定カメラで撮影した少ないフレーム数の画像をもとに, 任意の照明下における三次元物体を任意の姿勢に対して生成する手法を提案した. 線形結合係数の繰り返し計算を導入することにより, 三次元形状と画像基底の精度を向上することが可能となった. 特に, 初期特徴点が不十分な場合の精度向上が顕著であることから, テクスチャの少ない対象物体に対して有効な結果だと考えられる. また, 得られる形状を利用して, 任意の姿勢に対して任意光源下での画像を生成できることを実験的に示した.

ここでは, 完全拡散反射特性をもつ三次元物体表面を仮定しているため, 基底画像の数を 3 としている. 他の反射特性をもつ物体表面に対する拡張や影 (self-shadow), 隠蔽の扱い等が今後の課題である. 光源としては単一光源を仮定したが, 多光源への拡張は [6] により検討中である.

参考文献

- [1] P.N. Belhumeur, J.P. Hespanha, and Kriegman D.J. Eigenfaces vs. fisherfaces: recognition using class specific linear projection. *IEEE-PAMI*, Vol. 19, No. 7, pp. 711–720, 1997.
- [2] P.N. Belhumeur and D.J. Kriegman. What is the set of images of an object under all possible illumination conditions? *IJCV*, Vol. 28:3, pp. 245–260, 1998.
- [3] R. Epstein, P.W. Hallinan, and A. Yuille. 5 ± 2 Eigenimages suffice: an empirical investigation of low-dimensional lighting models. In *Proc. IEEE Workshop on Physics-based Modeling in CV*, 1995.
- [4] A.S. Georghiades, P.N. Belhumeur, and D.J. Kriegman. From few to many: Illumination cone models for face recognition under variable lighting and pose. *IEEE-PAMI*, Vol. 23, No. 6, pp. 643–659, June 2001.
- [5] A. Maki, M. Watanabe, and C.S. Wiles. Geotensity: Combining motion and lighting for 3d surface reconstruction. *IJCV*, Vol. 48, No. 2, pp. 75–90, 2002.
- [6] A. Maki and C. Wiles. Geotensity constraint for 3d surface reconstruction under multiple light sources. In *6th ECCV*, pp. 725–741, 2000.
- [7] J.L. Mundy and A. Zisserman, editors. *Geometric invariance in computer vision*. The MIT Press, 1992.
- [8] H. Murase and S.K. Nayer. 3d object recognition from appearance - parametric eigenspace method. *IEICE Transactions on Information and Systems (D-II)*, Vol. 77-D-II, No. 11, pp. 2197–2187, 1994.
- [9] A. Nakashima, A. Maki, and K. Fukui. Constructing illumination image basis from object motion. In *7th ECCV*, pp. 195–209. Springer-Verlag, 2002.
- [10] A. Shashua. *Geometry and photometry in 3D visual recognition*. PhD thesis, Dept. of Brain and Cognitive Science, MIT, 1992.
- [11] A. Shashua and T. Riklin-Raviv. The quotient image: Class-based re-rendering and recognition with varying illuminations. *IEEE-PAMI*, Vol. 23, No. 2, pp. 129–139, 2001.
- [12] D. Weinshall and C. Tomasi. Linear and incremental acquisition of invariant shape models from image sequences. In *4th ICCV*, pp. 675–682, 1993.
- [13] 中島朗子, 牧淳人, 福井和広. 運動物体からの照明変動画像基底の合成. 信学技報, Vol. PRMU2002-55, pp. 63–68, 2002.
- [14] 福井和広, 山口修, 鈴木薫, 前田賢一. 制約相互部分空間法を用いた環境変動にロバストな顔画像認識 - 照明変動の影響を抑える制約部分空間の学習 -. 信学論 (D-II), Vol. J82-D-II, No. 4, pp. 613–620, 1999.