

簡易ユーザ入力による景観画像の 3D シーン生成システム

飯塚里志[†] 金森由博[†] 三谷純[†] 福井幸男[†]

本稿では 1 枚の景観画像から 3D モデルを容易に生成できるシステムを提案する。生成される 3D モデルは背景モデルと前景物モデルから構成され、ユーザが地面領域と立体物との「境界線」をインタラクティブに指定することで、各モデルの 3 次元座標が算出される。さらに提案システムにより前景物モデリングやテクスチャ生成を簡単に行うことができ、簡素ながら十分な 3D 効果の得られる 3D シーンが生成されることを示す。

An efficient modeling system for 3-D scenes from a single landscape image

Satoshi Iizuka[†] Yoshihiro Kanamori[†] Jun Mitani[†]
and Yukio Fukui[†]

This paper presents a system for creating a 3D model easily from a single image. The created model is composed of a background and foreground objects whose coordinates are calculated based on a “boundary” between a ground region and objects. We show that the proposed system enables efficient modeling of foreground objects, easy creation of their textures, and rapid construction of a 3D scene which is simple but produces sufficient 3D effects.

1. はじめに

近年、フォトリアリスティックなシーンを効率よく生成するため、イメージベースによるレンダリング手法が多く研究されている。この手法は画像をテクスチャとして

用いることで複雑な物体形状を正確にモデル化しなくてもリアリティのある 3D 空間が作成でき、高速なレンダリングが可能となる。作成された空間の視点を移動することにより、ユーザはその空間を実際に歩いているような体験をすることができ、バーチャルツアーなど特定の景観を 3 次元的に表現するときによく用いられる。これらは Snively らが開発した *Microsoft photosynth*¹⁶⁾ や *Google maps* におけるストリートビュー⁶⁾ などの一般的なアプリケーションにも見られ、非常に身近な存在になってきている。しかし、そのような 3D シーンを作成するためには入力として多くの写真を必要とする。また、このような 3D シーンの作成は非常に手間のかかる作業である。

本稿では 1 枚の画像を入力としてその 3D シーンを容易に作成できるシステムを提案する。本システムではまずユーザが地面と立体物の境界線をインタラクティブに指定する。この境界線によってモデルの 3 次元座標が算出され、3D シーンモデルが生成される。また、入力画像の領域分割⁵⁾ とグラフカットベースの最適化¹³⁾ により簡単に前景物の抽出とそのモデリングを行うことができ、前景物を含む景観画像でも 3D シーンモデルが容易に作成できる。生成されるモデルは少数のポリゴンに画像テクスチャがマッピングされた単純なモデルであり、入力画像の 3 次元構造を正確に再現するものではないが、ユーザに十分な 3D 効果を与えることができる。

本手法は入力画像を 1 枚しか必要としない。そのため、インターネット上で手に入れた海外の風景写真や 1 つの視点からしか見ることのできない風景画などを入力として用いることで、通常 3 次元的に見ることができないような風景画像の 3D ウォークスルーが可能となる。もちろん提案手法では 3D モデルの生成がうまくいかないような画像も多く存在するが、それでも本システムは様々な景観画像の 3D シーンを簡単に作成でき、有用なツールとなることを示す。

2. 関連研究

2D 画像から 3D シーンを生成する手法は数多く研究されている。イメージベースのレンダリング手法を用いた一般的なアプリケーションである *Quicktime VR*⁴⁾ や *Microsoft photosynth*¹⁶⁾、*Google maps* のストリートビュー⁶⁾ などは、ステレオアルゴリズムなどを用いて広範囲にわたる 3D 空間を作成することができるが、非常に多くの写真と特殊な装置を必要とする。このような多くの画像を用いて 3D シーンを作成する手法に対し、1 枚の画像から 3D シーンを作成することはその情報量の少なさからより困難な問題である。

以下では 1 枚の画像から 3D モデルを生成する代表的な手法について述べる。Hoiem らは機械学習や領域分割を用いて、屋外写真の地面、空、垂直線の 3 つの領域を推定し、この情報から 3D シーンを自動で生成する手法を示した⁷⁾。また Saxena らは画像をスーパーピクセルと呼ばれる小さな領域に分け、それらのピクセル同士のグラデーション

[†]筑波大学大学院システム情報工学研究科コンピュータサイエンス専攻
Department of Computer Science, University of Tsukuba

ョンの変化などを分析することによって 3Dシーンを生成する手法を提案した¹⁴⁾。これらの手法は 3Dシーンを自動で生成できる反面、奥行き推定や領域推定に失敗する 경우가少なくなく、適用可能な画像に限られる。また、前景物のある入力画像には適用が困難である。これらの手法に対し、Ohらは 1 枚の画像を入力として、画像中の木や建物を階層的に並べることで 3Dシーンをインタラクティブに生成するシステムを提案した¹¹⁾。このシステムでは画像の照明なども編集でき、非常に良好な結果を得られるが、領域分割やその奥行き割り当てなどを手動で行う必要があり、1つの 3Dシーンを作成するのに非常に手間がかかる。また、KangらはTour Into the Pictureと呼ばれる手法⁸⁾をもとに、消失線にもとづいて 1 枚の景観画像から単純な 3D シーンモデルを生成する手法を提案した⁹⁾。この手法は単純なシーンモデルで十分な 3D 効果を提供できるが、地平線や水平線が存在するような広い景観画像を入力として想定しているため、建物の多い市街地や屋内のように複雑な形状が現れる景観画像には適用が困難である。また、モデルのテクスチャ生成を一般的なペイントツールを用いて個別に行う必要があり、これは手間のかかる作業である。

提案手法では、入力画像の立体物と地面領域の境界線から 3D モデルを生成する。これにより既存モデルよりも多くの景観画像に対応することが可能になる。また、既存手法では示されていない前景物の抽出やテクスチャ生成を簡単に行うことができるツールを提案し、3D シーンモデルが効率よく作成できることを示す。

3. 提案システム

提案システムでは地面領域をもつ 1 枚の景観画像を入力とし、その 3D シーンは 1つの背景モデルと複数の前景物モデルで構成される。これらのモデルは少数のポリゴンに入力画像の各領域がテクスチャマッピングされたものであり、背景モデルの地面領域に前景物モデルが垂直に設置される。このシステムにおいて、3D シーン作成のためにユーザは 2つの簡単な作業を行う。すなわち入力画像における境界線と前景物の指定である。これらの作業は入力画像上で行われ、3次元的な編集を必要としない。

3.1 境界線

入力画像の 3D モデルは境界線により決定される。ここで述べる境界線とは「地面」領域と「壁」領域の境目を指し、ユーザによって折れ線で指定される。例えば 図 1(a)のような写真では地面と建物の境目が指定される。ここでは建物や空が「壁」領域となる。また、図 7(上段)のように境界に曲線が現れるような画像に対しても、折れ線の頂点を増やした近似的な曲線で境界線を指定することができる。この境界線の各頂点座標にもとづいて背景モデルの 3次元構造が決定され、前景物モデルの 3次元位置を算出することができる(図 1(b))。さらに境界線は背景テクスチャ生成に必要な前景物領域の穴埋めの際に拘束条件としても用いられ(4.1 節で後述)、背景テクスチャ



(a) 入力画像

(b) 3D モデル

(c) 視点移動

図 1 提案システムによる 3D シーンの生成。入力画像の境界線が赤色、前景物が青色で示されている(a)。境界線から 3D モデル(b)が算出され、視点移動(c)が可能になる。

生成の精度を向上させる役割も持つ。

3.2 前景物

入力画像の各前景物はユーザによって指定される。本稿における前景物とは境界線によって分割された地面領域に置かれている立体物のことであり、基本的に 1 枚の板状ポリゴンでモデル化される。しかし前景物の抽出は通常手間のかかる作業である。そこで本システムでは、領域分割を用いた選択領域の最適化とグラフカットを用いた前景領域の最適化を併せた前景物抽出手法を提案する。この手法ではユーザははじめに前景物を粗く囲み、この情報をもとにシステムが前景物領域の最適化を行う。もし前景物が正確に抽出できていない場合はその領域をユーザが指定し、それをシステムが学習してもう一度最適化を行う。この手法により少ないユーザ入力で簡単に前景物を抽出することが可能となる。さらに、前景物領域の中で地面に接している部分をユーザが明示的に指定することで遠近感をもった前景物をモデル化できることを示す。

4. シーンモデリング

この節ではシステムの基盤となるアルゴリズムについて述べる。最初に、境界線にもとづく背景モデルの生成アルゴリズムとそのテクスチャ生成について述べる。次に簡易ユーザ入力による前景物の抽出手法とそのモデリングについて説明する。さらに、前景物モデルに適用されるビルボード変換や接地制約について述べ、これによりシーンのリアリティが向上することを示す。

4.1 背景モデル

入力画像はユーザによって指定された折れ線にもとづいて地面ポリゴンと複数の壁ポリゴンに分割される(図 2)。この各頂点に適切な座標を割り当てることで背景モデルが構築される。スクリーン座標系において左下を原点とし右方向を $+x$ 、上方向

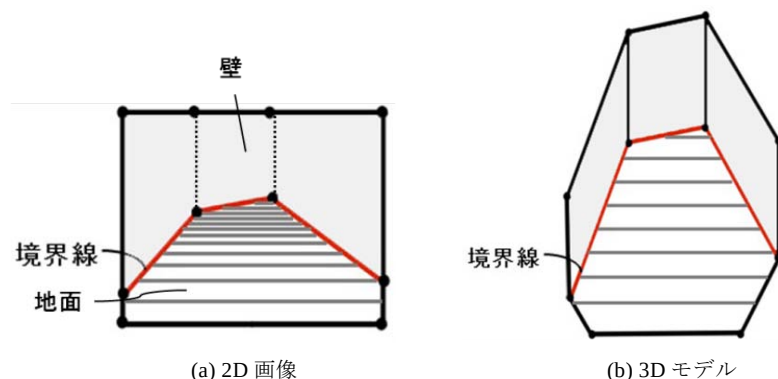


図 2 2D 入力画像(a)とその 3D モデル(b). 境界線によって入力画像は地面領域と壁領域に分割され、境界線の頂点座標から 3D モデルが生成される。

を $+y$ とする。また、3D モデルにおいて原点はカメラ位置と一致し、視線方向を $+z$ 、カメラの焦点距離を f とする。ここで同次座標 (x, y, z, w) は $w \in [0, 1]$ が小さいほど 3 次元空間の奥に透視投影変換される。よって入力画像のスクリーン座標の原点 $\mathbf{P}_0$ を (x_0, y_0) 、境界線上で最も y 座標が大きな頂点 $\mathbf{P}_M$ を (x_M, y_M) とし、ワールド座標系における $\mathbf{P}_0, \mathbf{P}_M$ をそれぞれ $(x'_0, y'_0, f), (x'_M, y'_M, f)$ とすると、その同次座標 $\mathbf{P}'_0, \mathbf{P}'_M$ は以下のように表される。

$$\mathbf{P}'_0: (x'_0, y'_0, f, 1) \quad \mathbf{P}'_M: (x'_M, y'_M, f, w_{min})$$

ここで w_{min} とは小さな正の値であり、本システムでは 0.2 としている。この 2 点を基準として入力画像の地面領域の各頂点 $\mathbf{P}_i(x_i, y_i)$ が以下のように算出される。

$$\mathbf{P}'_i: (x'_i, y'_i, f, w_i)$$

ただし、

$$w_i = \frac{y_i - y_0}{y_M - y_0} w_{min} + 1 - \frac{y_i - y_0}{y_M - y_0} \quad (1)$$

また、残りの壁領域の各頂点は壁と地面は垂直であるという制約条件のもとに算出される。これにより視線方向に奥行きをもつ 3D モデルが生成される。

生成されるモデルはKangらの消失線の理論⁹⁾にもとづいているが、本手法は各頂点の y 値によって座標を決定しているため、Kangらが示した算出手法よりも単純に計算でき、さらにこの方程式を用いて前景物モデルの同次座標も算出できる。

背景テクスチャ 生成されたポリゴンモデルに背景テクスチャをマッピングすることで背景モデルが完成する。本稿における背景とは入力画像の中で前景物が含まれな



(a) 入力画像 (b) 補完結果 (制約なし) (c) 補完結果 (制約あり)

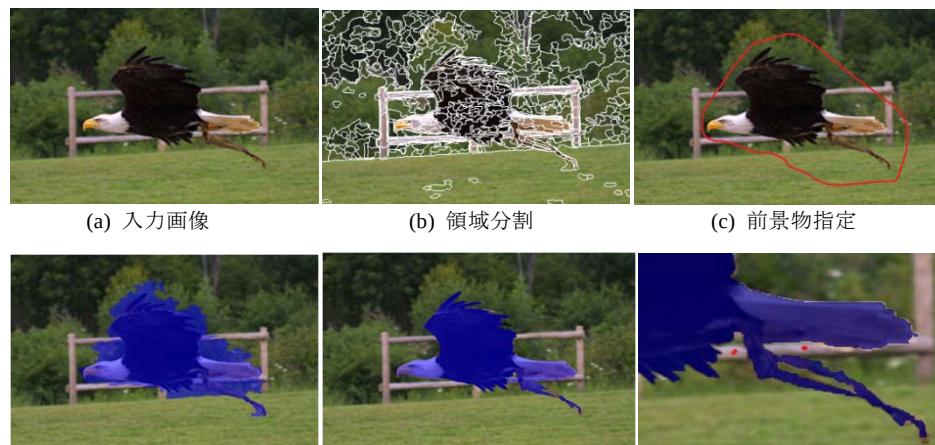
図 3 背景テクスチャの生成。入力画像(a)では境界線が赤色、前景物が青色で示されている。制約なしで補完を行うと不自然なテクスチャ(b)が生成されるが、境界線を制約条件として補完を行うと良好な結果(c)が得られる。

い領域である。しかし前景物領域を取り除いた入力画像を背景テクスチャとして用いると前景物の部分が「穴」になってしまう。これを避けるため、前景物領域はその他の背景領域で補完する必要がある。本システムでは補完のためにBarnesらのPatchMatchと呼ばれる手法²⁾を用いる。彼らの手法により、計算コストの高い類似領域探索¹⁵⁾を高速化し、画像を高速に自動補完することができる。本システムではこの手法を実装し、指定された前景物が抜き出されると同時にその領域が自動的に補完される。よって、すべての前景物の抽出が完了すると同時に完全な背景テクスチャが生成される。しかし画像補完は対象領域が大きくなるほど精度が低くなってしまう(図 3(b))。Barnesらは、画像補完の際にガイド線上の類似領域探索をその線上に拘束することで画像構造を補完に反映させる手法を提案した。本システムではこの理論にもとづき、境界線を補完の拘束条件として用いることで前景物の補完精度を向上させている(図 3(c))。つまり、境界線は 3 次元座標の算出と背景テクスチャ生成の精度向上という 2 つの役割を果たしている。

4.2 前景物モデル

前景物の抽出は画像編集の際に最も手間のかかる作業の 1 つである。既存手法では前景をブラシで直接塗りつぶしていくペイントベースの抽出手法¹¹⁾¹²⁾や背景と前景の境界をなぞっていく境界線ベースの抽出手法¹⁰⁾が示されている。しかし、これらは対象となる前景物の形状によっては細かく正確な作業が必要となる。提案システムではユーザ入力となるべく少なくし、かつ細かく正確な作業を必要としない前景物抽出を目指し、領域分割とグラフカットベースの手法を組み合わせることで前景最適化を行う。

本システムでは入力画像は読み込まれると同時に領域分割が適用される(図 4(b))。この領域分割のために、我々はComaniciuらの手法⁵⁾を用いる。この手法では入力画像の色を表す $L \cdot u \cdot v$ と位置を表す x, y の 5 属性を特徴としてMean Shiftを行い、近接

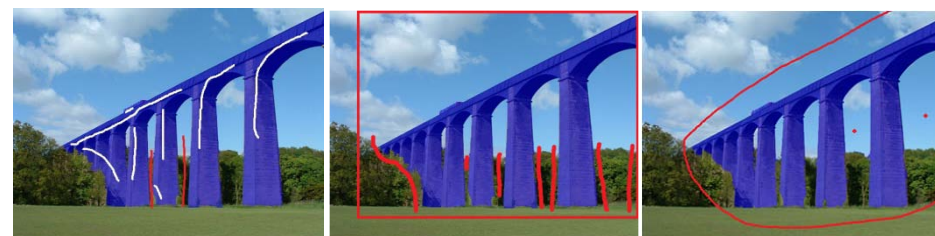


(a) 入力画像 (b) 領域分割 (c) 前景物指定
(d) 領域情報を用いた最適化 (e) GrabCutによる最適化 (f) 追加ユーザ入力による修正
図 4 前景物の抽出. 入力画像(a)は前処理として領域分割されている(b). ユーザが前景物を粗く囲む(c)と, そのガイド線の外側とガイド線が含まれている領域が背景として処理される(d). この情報をもとにGrabCut¹³⁾によって前景領域が最適される(e). もしうまく抽出できていない場合はユーザがその部分を指定して修正することができる(f).

領域同士の色空間におけるユークリッド距離が閾値以下のものを統合することで領域分割を行う. この手法は多くの画像で高い精度を実現することが示されている.

前景物抽出のため, まずユーザはなげなわツールのようなインタフェースを用いて前景物領域を粗く囲む(図 4(a)). 囲まれた領域の外側は背景領域とし, さらにガイド線が含まれる各領域は背景であると推測できるため, その領域もすべて背景とみなす(図 4(d)). これを初期状態とし, さらにRotherらのGrabCutと呼ばれる手法¹³⁾にもとづいて領域の最適化を行う(図 4(e)). 通常, GrabCutはグラフカットによる分割結果から前景と背景の色分布を再学習し, 反復的に前景分割を行うことでその精度を向上させている. しかし, 反復処理は収束まで時間がかかりインタラクティブな編集には向かないことが多い. 本システムでは, 先に領域分割を用いた領域の最適化を行っておくことで繰り返し処理を行わずに 1 度だけのGrabCutによる最適化で十分な前景物分割が可能となる.

しかし正確な前景物抽出は困難な課題であり, 上記の処理だけではうまく抽出されない場合がある. GrabCut ではユーザが明示的に前景や背景領域を指定し, この情報をもとに再度分割を行うことでより正確な前景抽出を行う編集機能が示されている. 本システムではこれに上記の領域分割を用いた最適化を追加し, ユーザによる前景/



(a) Photoshop Quick Selection (b) GrabCut (c) 提案手法

図 5 前景物の抽出結果. 赤線が背景のガイド線, 白線が前景のガイド線を表している. Photoshop CS5 の Quick Selection(a)や Rother らの GrabCut(b)に比べ, 提案手法(c)は粗く少ないユーザ入力で前景物を抽出できている.

背景指定, 領域分割を用いた最適化, グラフカットによる最適化の 3 段階を繰り返すことでより粗く少ないユーザ入力で正確な前景物抽出を行うことができる.

本システムによる前景物抽出結果を図 5 に示す. 入力画像のサイズは 800×600 ピクセルである. 提案手法のユーザ入力に対する最適化処理時間は約 0.41 秒であった. また, ユーザの作業時間も含めて前景物抽出にかかった時間は, 強力な画像編集ツールである Adobe Photoshop CS5 の Quick Selection ツールでは 107 秒, 我々の実装による Rother らの GrabCut では 115 秒, 提案手法では 18 秒であった. よって図 5 に示す例では, 提案手法は他の手法と比べ, 前景抽出にかかる時間が約 85%低減されている. ただし, 本手法の前景抽出精度は GrabCut による最適化に依存しており, 既存手法よりも精度が向上するわけではない. しかし他の手法に比べてより粗く少ないユーザ入力で前景物を抜き出すことができ, ユーザの負担を軽減することが可能となる.

提案手法における前景物抽出は, 前処理で行われる領域分割において背景と前景物が同じ領域として分割された場合にはうまく最適化が行えない. しかし本システムでは領域分割を用いた最適化を行わず, ブラシによる直接塗りつぶしと GrabCut による最適化だけを用いた前景抽出や, 単純にブラシによる塗りつぶしだけで前景物を抽出することもできるため, ほぼすべての前景物に対応できる.

このようにして抽出された前景物画像を 1 枚の四角形ポリゴンに前景物テクスチャとしてマッピングすることで前景物モデルが生成され, 背景モデルの地面領域に垂直に配置される. この前景物モデルの各頂点の 3 次元座標は, 前景物領域の中で最も小さい y 座標にもとづき式(1)と垂直条件から計算できる.

ビルボード変換 前景物モデルは板状ポリゴンであるため, 視点を横へ移動すると前景物の立体感が失われてしまう. これを解決するため, 本システムではビルボード変換を前景物モデルに適用する. ビルボード変換とは対象のポリゴンが常に視点方向を

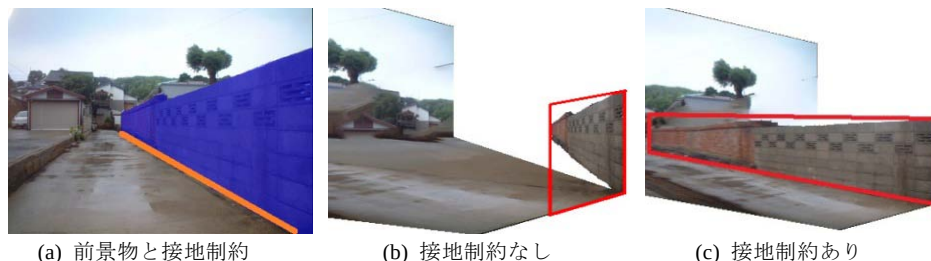


図 6 前景物の接地制約. 青色で示されている前景物に橙色の直線で示される接地条件を加えることで(a), 前景物に正確な 3 次元位置を与えることができる(c). (b)(c)の赤線は前景物のポリゴンモデルを表している.

向くように座標変換を施す手法であり, 本システムでは木や円筒状の前景物に適用される. このビルボード変換行列 Π_b は以下のように表される.

$$\Pi_b = T^{-1}RT$$

ただし,

$$T = \begin{pmatrix} E & V \\ 0^t & 1 \end{pmatrix}, \quad R = \begin{pmatrix} R_y & 0 \\ 0^t & 1 \end{pmatrix}$$

この式において V は視点座標と対象ポリゴンの回転軸座標の差を表すベクトルであり, R_y は y 軸周りの回転行列である. これにより, 対象の前景物ポリゴンが常に視点方向を向くようになり, 前景物の立体感を表現することができる.

接地制約 四角形の板状ポリゴンである前景物モデルの各頂点には同じ奥行きが与えられている. しかし前景物が奥に向かって地面領域に置かれている場合には不自然なモデルが生成されてしまう. そこで本システムでは前景物が地面領域に接している部分を直線で指定することで, 前景物モデルの 3 次元座標を修正できる機能を実装している (図 6). この機能を用いることで前景物に立体形状を与えることも可能となる. このような前景物モデルは回転すると不自然なため, ビルボード変換は適用しない.

4.3 並列処理

上記で述べた画像補完や前景物最適化処理の高速化はインタラクティブな編集には重要な問題である. Bernsらは画像をタイル状に分割して並列に処理することでより高速に類似領域探索が行えることを示した³⁾. 本システムではこの並列処理手法を用いて, 画像の補完を複数のCPUコアによって並列処理している. これによりほぼコア数に比例して類似領域探索処理が高速化される. また, *GrabCut*処理の際に計算コストの高いグラフ生成処理を並列処理することで前景物抽出の処理性能を向上させている. これは 2 コアで約 20%の処理速度の向上がみられた.

5. 結果と考察

本システムはライブラリとして OpenGL, GLUT, GLUI, OpenMP を用いて C++ 言語で実装し, Intel Core i7 620M (2.67GHz, 4.00GB RAM) と NVIDIA Quadro NVS 3100M グラフィックカードが搭載された PC 上で実行した. 使用した画像のサイズは全て 0.5 から 1.0 メガピクセルの範囲内である.

図 1 では 2 基の街灯が前景物として指定され, 背景の境界が 5 個の頂点をもつ折れ線で指定されている. この街灯はビルボード変換が適用され, 横からの視点に対しても立体感のある自然な前景物が生成される. 街灯のような円柱状の物体や一般的な木などの前景物は, 回転してもあまり形状が変わらない場合が多く, ビルボード変換で十分に対応することができる. しかし, 本手法では地面や壁のそれぞれの領域は 1 つの平面としてモデル化されているため, 視点によっては階段が後ろの建物に張り付いているような違和感をユーザに与えることがある.

図 7 (上段) は曲線と直線を含む境界線をもつ屋内の写真である. このように前景物がない画像は境界線を指定するだけで 3D シーンを生成することができる. また, 図 7 (中段) では建物に折れ線で接地制約を与えることで, 建物が 2 枚のポリゴンで立体的にモデル化されている. このように接地制約を折れ線で指定することにより, 通常 1 枚の板状ポリゴンモデルである前景物を, 複数のポリゴンから構成される立体的なモデルにすることができる. さらに図 7 (下段) のように入力が絵画であっても提案システムによりその 3D シーンを作成することができる.

図 8 は Hoiem らの手法⁷⁾と提案手法の比較である. Hoiem らのモデルでは人が地面や建物と一体化してしまっているのに対し, 提案システムのモデルは人が前景物としてモデル化され, その背景もきれいに補完されているため, 自然なシーンが生成されている.

これらの 3D シーンを作成するのににかかった時間はすべて 3 分以下であった. この作業時間において大きな割合を占めるのは, 入力画像の前景物抽出にかかる時間である. 例えば前景物がない図 7 (上段) の場合, 3D シーンを作成するのににかかった時間が 14 秒であったのに対し, 3 つの前景物をもつ図 7 (中段) では作業時間は 127 秒であった. このように 3D シーンを作成するための作業時間は, 前景物の数に応じて増加していく. よって, 提案システムにおける前景物抽出の高速化は 3D シーンを作成する上で非常に重要な役割を果たしていることがわかる.

6. 結論

本稿では 1 枚の景観画像から容易に 3D シーンを生成できるシステムを提案した. 生成されるシーンモデルは背景モデルと前景物モデルから構成され, ユーザがインタラクティブに指定した地面と壁の境界線からその 3 次元座標が算出される. この境界



(a) 入力画像 (b) 境界線と前景物の指定 (c) 視点移動

図 7 3D シーンの生成. 入力画像(a)の境界線が赤線, 前景物が青色, 接地制約が橙色の線で示されている(b). (c)では生成されたシーンの視点を移動している.



(a) 入力画像 (b) Hoiem らの手法 (c) 提案システム

図 8 既存手法との比較. Hoiemらの手法⁷⁾に比べ, 提案システムは自然な 3Dシーンが生成されている.

線を用いたモデリングにより, 様々な景観画像の 3D シーンを作成することが可能となる. さらに提案システムでは前景物の簡易抽出やモデリング, 背景テクスチャの自動生成を行うことができる. これによりユーザの負担が少ない効率的なシーンモデリングが可能となる.

今後の研究では, ユーザが前景物に対して立体のプリミティブを当てはめたり, 立体形状を表す稜線をひいたりすることで, それを反映した立体的な前景物モデルが作成できる機能を実現したいと考えている. また, 透視投影によって引き伸ばされるテクスチャを高解像度化することでよりシーンのリアリティを向上させたい.

参考文献

- 1) Adobe Photoshop. <http://www.adobe.com/support/photoshop/>.
- 2) Barnes et al., PatchMatch: a randomized correspondence algorithm for structural image editing. In Proc. of ACM SIGGRAPH 2009, 24.
- 3) Barnes et al., The Generalized PatchMatch Correspondence Algorithm. In Proc. of ECCV 2010, 29-43.
- 4) Chen, Quicktime VR - An Image-based Approach to Virtual Environment Navigation. In Proc. of ACM SIGGRAPH 1995, 29-38.
- 5) Comaniciu and P.Meer. A robust approach toward feature space analysis. In Proc. of TPAMI 2002, 24,5.
- 6) Google maps. <http://maps.google.com/>
- 7) Hoiem et al., Automatic photo pop-up. In Proc. of ACM SIGGRAPH 2005. 24, 3, 577-584.
- 8) Horry et al, Tour Into the Picture: Using a Spidery Mesh Interface to Make Animation from a Single Image. In Proc. of ACM SIGGRAPH 1997, 225-232.
- 9) Kang et al., Tour Into the Picture using a Vanishing Line and its Extension to Panoramic Images. In Proc. of EuroGraphics 2001, 132-141.
- 10) Mortensen et al., Tobogan-based intelligent scissors with a four parameter edge model. In Proc. of CVPR 1999, vol. 2, 452-458.
- 11) Oh et al., Image-based modeling and photo editing. In Proc. of ACM SIGGRAPH 2001, 433-442.
- 12) Olsen et al., Edge-respecting brushes. In Proc. of UIST 2008, 171-180.
- 13) Rother et al., GrabCut: Interactive Foreground Extraction Using Iterated Graph Cuts, In Proc. of ACM SIGGRAPH 2004, 23, 3, 309-314.
- 14) Saxena et al., Make3D: Learning 3D scene structure from a single still image. In Proc. of TPAMI 2009, 824-840.
- 15) Simakov et al., Summarizing visual data using bidirectional similarity. In Proc. of CVPR 2008, 1-8.
- 16) Snavely et al., Photo tourism: Exploring photo collections in 3D. In Proc. of SIGGRAPH 2006, 835-846.