

Light Field Appearance Manifolds

Chris Mario Christoudias, Louis-Philippe Morency, and Trevor Darrell

Computer Science and Artificial Intelligence Laboratory
Massachusetts Institute of Technology
Cambridge, MA 02139 USA
{cmch, lmorency, trevor}@csail.mit.edu
<http://www.ai.mit.edu/projects/vip/>

Abstract. Statistical shape and texture appearance models are powerful image representations, but previously had been restricted to 2D or 3D shapes with smooth surfaces and lambertian reflectance. In this paper we present a novel 3D appearance model using image-based rendering techniques, which can represent complex lighting conditions, structures, and surfaces. We construct a light field manifold capturing the multi-view appearance of an object class and extend the direct search algorithm of Cootes and Taylor to match new light fields or 2D images of an object to a point on this manifold. When matching to a 2D image the reconstructed light field can be used to render unseen views of the object. Our technique differs from previous view-based active appearance models in that model coefficients between views are explicitly linked, and that we do not model any pose variation within the shape model at a single view. It overcomes the limitations of polygonal based appearance models and uses light fields that are acquired in real-time.

1 Introduction

Appearance models are a natural and powerful way of describing objects of the same class. Multidimensional morphable models [13], active appearance models [6], and their extensions have been applied to model a wide range of object appearance. The majority of these approaches represent objects in 2D and model view change by morphing between the different views of an object. Modelling a wide range of viewpoints in a single 2D appearance model is possible, but requires non-linear search [19]. Additionally, object self-occlusion introduces holes and folds in the synthesized target view which are difficult to overcome. Large pose variation is easily modelled using 3D; a polygonal 3D appearance model was proposed by Blanz and Vetter [3]. With their approach the view is an external parameter of the model and does not need to be modelled as shape variation. However, this technique is based on a textured polygonal mesh which has difficulty representing fine structure, complex lighting conditions and non-lambertian surfaces. Due to the accuracy of the 3D surfaces needed with their approach, the face scans of each prototype subject cannot be captured in real-time and fine structure such as hair cannot be acquired.

In this paper we propose a 3D active appearance model using image-based

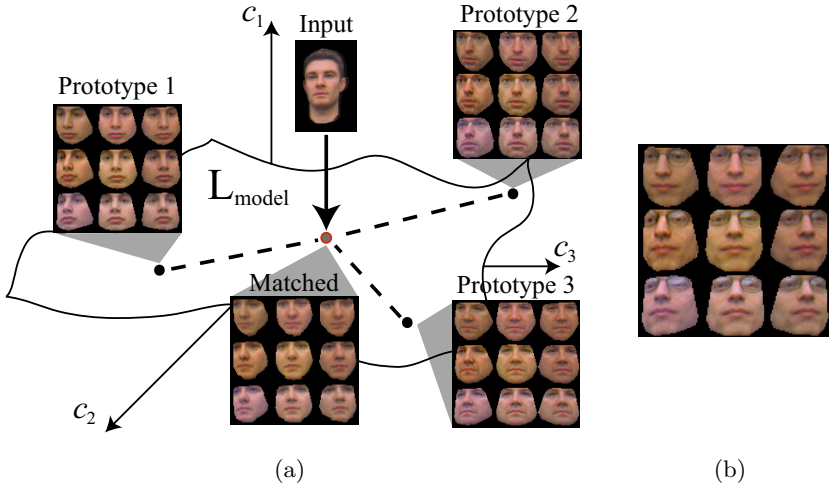


Fig. 1. (a) A light field appearance manifold L_{model} . Each point on the manifold is a 4D light field representing the 3D shape and surface reflectance of an object. The light field of an object is constructed by computing its projection onto the shape-texture appearance manifold. A 2D input image is matched to a point on this manifold by interpolating the shape and texture of neighboring prototype light fields. (b) A light field can capture non-lambertian effects (e.g. glasses).

rendering [14,11] rather than rendering with a polygonal mesh. We use a light field representation, which does not require any depth information to render novel views of the scene. With light field rendering, each model prototype consists of a set of sample views of the plenoptic function [1]. Shape is defined for each prototype and a combined texture-shape PCA space computed. The resulting appearance manifold (see Figure 1(a)) can be matched to a light field or 2D image of a novel object by searching over the combined texture-shape parameters on the manifold. We extend the direct search matching algorithm of [6] to light fields. Specifically, we construct a Jacobian matrix consisting of intensity gradient light fields. A 2D image is matched by rendering the Jacobian at the estimated object pose. Our approach can easily model complex scenes, lighting effects, and can be captured in real-time using camera arrays [23,22].

2 Previous Work

Statistical models based on linear manifolds of shape and/or texture variation have been widely applied to the modelling, tracking, and recognition of objects [2, 8,13,17]. In these methods small amounts of pose change are typically modeled implicitly as part of shape variation on the linear manifold. For representing objects with large amounts of rotation, nonlinear models have been proposed, but are complex to optimize [19]. An alternative approach to capturing pose variation

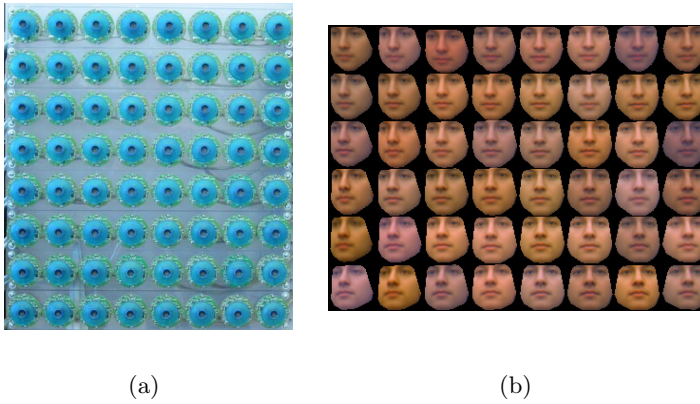


Fig. 2. (a) Light field camera array [23]. (b) A 6x8 light field of the average head. The light field prototypes were acquired using the 6 top rows of the camera array due to field of view constraints.

is to use an explicit multi-view representation which builds a PCA model at several viewpoints. This approach has been used for pure intensity models [16] as well as shape and texture models [7]. A model of inter-view variation can be recovered using the approach in [7], and missing views could be reconstructed. However, in this approach pose change is encoded as shape variation, in contrast to 3D approaches where pose is an external parameter. Additionally, views were relatively sparse, and individual features were not matched across views.

Shape models with 3D features have the advantage that viewpoint change can be explicitly optimized while matching or rendering the model. Blanz and Vetter [3] showed how a morphable model could be created from 3D range scans of human heads. This approach represented objects as simply textured 3D shapes, and relied on high-resolution range scanners to construct a model; non-lambertian and dynamic effects are difficult to capture using this framework. With some manual intervention, 3D models can be learned directly from monocular video [9,18]; an automatic method for computing a 3D morphable model from video was shown in [4]. These methods all used textured polygonal mesh models for representing and rendering shape.

Multi-view 2D [7] and textured polygonal 3D [3,9,18] appearance models cannot model objects with complex surface reflectance. Image-based models have become popular in computer graphics recently and can capture these phenomenon; with an image-based model, 3D object appearance is captured in a set of sampled views or ray bundles. Light field [14] and lumigraph [11] rendering techniques create new images by resampling the set of stored rays that represent an object. Most recently the unstructured lumigraph [5] was proposed, and generalized the light field/lumigraph representation to handle arbitrary camera placement and geometric proxies.

Recently, Gross et. al. [12] have proposed *eigen light fields*, a PCA-based appearance model built using light fields. They extend the approach of Turk and Pentland [21] to light fields and define a robust pose-invariant face recognition algorithm using the resulting model. A method to morph two lightfields was presented in [24]; this algorithm extended the classic Beier and Neely algorithm to work directly on the sampled lightfield representation and to account for self-occlusion across views. Features were manually defined, and only a morph between two (synthetically rendered) light fields was shown in their work.

In this paper we develop the concept of a *light field active appearance model*, in which 3 or more light fields are “vectorized” (in the sense of [2]) and placed in correspondence. We construct a light field morphable model of facial appearance from real images, and show how that model can be automatically matched to single static intensity images with non-lambertian effects (e.g. glasses). Our model differs from the multi-view appearance model of [7] in that we build a 4D representation of appearance with light fields. With our method, model coefficients between views are explicitly linked and we do not model any pose variation within the shape model at a single view. We are therefore able to model self-occlusion and complex lighting effects better than a multi-view AAM. We support this claim in our experimental results section.

3 Light Field Shape and Texture

In this section we provide a formal description of the shape and texture of a set of light field prototypes that define the appearance manifold of an object class. Let $L(u, v, s, t)$ be a light field consisting of a set of sample views of the scene, parameterized by view indices (u, v) and scene radiance indices (s, t) , and let $L_1, \dots, L_n$ be a set of prototype light fields with shape $X_1, \dots, X_n$. In general, for any image-based rendering technique, X_i is a set of 3D feature points which outline the shape of the imaged object. With a light field, no 3D shape information is needed to render a novel view of the object. It is therefore sufficient to represent the shape of each light field as the set of 2D feature points, which are the projections of the 3D features into each view. More formally, we define the shape, X , of a light field L as

$$X = \{x_{(u,v)} \mid (u,v) \in L\} \quad (1)$$

where $x_{(u,v)}$ is the shape in a view (u, v) of L . If the camera array is strongly calibrated its sufficient to find correspondences in two views and re-project to the remaining views. With only weak calibration and the assumption of a densely sampled array, feature points may be specified in select views of the light field and tracked into all other views.

Once shape is defined for each prototype light field, Procrustes analysis [10] is performed to place the shape of each object into a common coordinate frame. Effectively, Procrustes analysis applies a rigid body transformation to the shape of each light field such that each object is aligned to the same 3D pose. From

the set of normalized shapes X_i of each prototype, the reference shape X_{ref} is computed as

$$X_{ref} = \mathbf{M}_\alpha \bar{X} \quad (2)$$

where $\bar{X}$ is the mean shape of the aligned shapes and $\mathbf{M}_\alpha$ is a matrix which scales and translates the mean shape such that it is expressed in pixel coordinates (i.e. with respect to the height and width of each discrete view of a light field). The matrix $\mathbf{M}_\alpha$ constrains the shape in each view of the reference light field to be within the height and width of the view.

As in [2], the texture of a prototype light field is its “shape free” equivalent. It is found by warping each light field to the reference shape X_{ref} . As will be shown in the next section, this allows for the definition of a texture vector space that is decoupled from shape variation. Specifically, the texture of a light field L is defined as

$$G(u, v, s, t) = L(D(u, v, s, t)) = L \circ D(u, v, s, t) \quad (3)$$

where D is the mapping,

$$D : \mathcal{R}^4 \longrightarrow \mathcal{R}^4 \quad (4)$$

that specifies for each ray in L a corresponding ray in the reference light field L_{ref} and is computed using the shape of L and X_{ref} . Equation (3) may be thought of as a light field warping operation, a concept introduced by Zhang et al. [24]. As in [6], the texture of each prototype, G_i , is normalized to be under the same global illumination.

4 Light Field Appearance Manifolds

As illustrated in the previous section, once a reference is defined, each prototype light field may be described in terms of its shape and texture. The linear combination of texture and shape form an appearance manifold: given a set of light fields of the same object class, the linear combination of their texture warped by a linear combination of their shape describes a new object whose shape and texture are spanned by that of the prototype light fields. Compact and efficient linear models of shape and texture variation may be obtained using PCA, as shown in [6]. Given the set of prototype light fields $L_1, \dots, L_n$, each having shape X_i and texture G_i , PCA is applied independently to the normalized shape and texture vectors, X_i and G_i to give

$$\begin{aligned} X &= \bar{X} + \mathbf{P}_s \mathbf{b}_s \\ G &= \bar{G} + \mathbf{P}_g \mathbf{b}_g \end{aligned} \quad (5)$$

Using Equation (5), the shape and texture of each model light field is described by its corresponding shape and texture parameters $\mathbf{b}_s$ and $\mathbf{b}_g$. As there may exist a correlation between texture and shape, a more compact model of shape and texture variation is obtained by performing a PCA on the concatenated

shape and texture parameter vectors of each prototype light field. This results in a combined texture-shape PCA space:

$$\begin{aligned} X &= \bar{X} + \mathbf{Q}_s \mathbf{c} \\ G &= \bar{G} + \mathbf{Q}_g \mathbf{c} \end{aligned} \quad (6)$$

where as in [6],

$$\begin{aligned} \mathbf{Q}_s &= \mathbf{P}_s \mathbf{W}_s^{-1} \mathbf{P}_{cs} \\ \mathbf{Q}_g &= \mathbf{P}_g \mathbf{P}_{cg} \end{aligned} \quad (7)$$

and $\mathbf{W}_s$ is a matrix which comensurates the variation in shape and texture when performing the combined texture-shape PCA. In our experiments we use $\mathbf{W}_s = r\mathbf{I}$ where $r = \sqrt{\sigma_s^2/\sigma_g^2}$. Here σ_s^2 and σ_g^2 represent the total variance of the normalized shape and texture. Equation (6) maps each model light field to a vector $\mathbf{c}$ in the combined texture-shape PCA space. To generalize the model to allow for arbitrary 3D pose and global illumination, Equation (6) may be re-defined as follows,

$$\begin{aligned} X_m &= S_t(\bar{X} + \mathbf{Q}_s \mathbf{c}) \\ G_m &= T_u(\bar{G} + \mathbf{Q}_g \mathbf{c}) \end{aligned} \quad (8)$$

where S_t is a function that applies a rigid body transformation to the model shape according to a pose parameter vector $\mathbf{t}$, T_u is a function which scales and shifts the model texture using an illumination parameter vector $\mathbf{u}$, and the parameter vectors $\mathbf{t}$ and $\mathbf{u}$ are as defined in [6]. Note, the reference light field has parameters $\mathbf{c} = 0$, $\mathbf{t} = \alpha$ and $\mathbf{u} = 0$, where α is a pose vector that is equivalent to the matrix $\mathbf{M}_\alpha$ in Equation (2).

The light field appearance manifold is defined as,

$$L_{model} = G_m \circ D_m \quad (9)$$

where L_{model} is a model light field that maps to a point on the appearance manifold and D_m is a 4D deformation field which maps each ray in the reference light field to a ray in the model light field and is computed using the shape of the model light field, X_m , and the shape of the reference light field, X_{ref} . Note, Equation (9) suggests that an optical flow technique may also be used to represent shape as in [13] to build a light field active appearance model. We have implemented both approaches, and below report results using the feature-based shape representation of Section 3.

5 Model Matching

In this section, we show how to generalize the matching technique of [6] to light fields. We first illustrate how to match a light field and then discuss the more interesting task of fitting a model light field to a single 2D image.

Matching to a Light Field. A novel light field, L_s , is matched to a point $\tilde{\mathbf{c}}$ on the texture-shape appearance manifold by minimizing the following non-linear objective function:

$$E(\mathbf{p}) = |G_m - G_s|^2 \quad (10)$$

where $\mathbf{p}^T = (\mathbf{c}^T | \mathbf{t}^T | \mathbf{u}^T)$ are the parameters of the model, G_m is the model texture and G_s is the normalized texture of L_s assuming it has shape X_m . G_s is computed by warping L_s from X_m to the reference shape X_{ref} . The model shape and texture are computed at $\mathbf{p}$ using Equation (8).

The direct search gradient descent algorithm of [6] is easily extendible to a light field active appearance model. In [6] a linear relationship for the change in image intensity with respect to the change in model parameters was derived via a first order Taylor expansion of the residual function $\mathbf{r}(\mathbf{p}) = G_m - G_s = \delta \mathbf{g}$. In particular, given a point $\mathbf{p}$ on the manifold, the parameter gradient that minimizes the objective function (10) was computed as, $\delta \mathbf{p} = -\mathbf{R} \delta \mathbf{g}$, where the matrix $\mathbf{R}$ is the pseudo-inverse of the Jacobian, $\mathbf{J} = \frac{\partial \mathbf{r}}{\partial \mathbf{p}}$, derived from the Taylor expansion of the residual function.

In a 2D active appearance model the columns of the Jacobian are intensity gradient images which model how image intensity changes with respect to each model parameter and vice versa. Analogously, the Jacobian of a light field active appearance model represents the change in light field intensity with respect to the change in model parameters, each of columns representing light field intensity gradients that describe the intensity change across all the views of a light field. Consequently, the algorithm for minimizing Equation (10) follows directly from [6]. As in a 2D AAM, the Jacobian is learned via numerical differentiation.

Matching to an Image. A more interesting extension of the AAM framework arises when performing direct search to match a light field AAM to a single 2D image; with a light field the Jacobian matrix is rendered based on pose. A novel image I_s is matched to a point on the light field appearance manifold by minimizing the objective,

$$E(\mathbf{p}, \epsilon) = |F(G_m, \epsilon) - g_s|^2 \quad (11)$$

where ϵ is the camera pose of I_s , F is a function that renders the pose ϵ of the model texture [14,5] and g_s is the texture of I_s assuming it has shape x_m . g_s is computed by warping I_s from x_m to the reference shape x_{ref} . Both 2D shapes are obtained by rendering X_m and X_{ref} into view ϵ using,

$$x = F_x(X, \epsilon) \quad (12)$$

where F_x is a variant of the light field rendering function F : it renders shape in view ϵ via a linear interpolation of the 2D shape features defined in each view of X .

Overall, the objective function in Equation (11) compares the novel 2D image to the corresponding view in L_{model} . Minimizing this objective function fits a model light field, L_{model} , that best approximates I in view ϵ . An efficient way to

optimize Equation (11) is by defining a two step iteration process, in which the pose ϵ is optimized independently of the model parameters $\mathbf{p}$. The pose ϵ may be computed via an exhaustive search of the average light field, L_{ref} , in which cross-correlation is used to initialize ϵ to a nearby discrete view of the model light field. The pose parameter $\mathbf{t}$ is used to further refine this pose estimate during matching.

Once ϵ is approximated, direct search may be employed to match I to a point on the texture-shape appearance manifold. As previously discussed, each column of the Jacobian, $\mathbf{J}$ of a light field active appearance model is a light field intensity gradient. To approximate the intensity gradient in view ϵ of the target image I , light field rendering is applied to each column of $\mathbf{J}$. This yields a “rendered” Jacobian matrix, $\mathbf{J}_\epsilon$, specified as,

$$\mathbf{J}_\epsilon^i = F(\mathbf{J}^i, \epsilon), i = 1, \dots, m \quad (13)$$

where $\mathbf{J}^i$ represents column i of the matrix $\mathbf{J}$ and m is the number of columns in $\mathbf{J}$. Note similar to the model and image textures of Equation (10) the columns of $\mathbf{J}_\epsilon$ have shape x_{ref} defined above.

Using $\mathbf{J}_\epsilon$, optimizing Equation (11) is analogous to matching I to a 2D AAM. Thus, as in Equation (10), the direct search gradient descent algorithm of [6] is used to minimize Equation (11), with one exception. In [6] the normalized mean of the texture vectors is used to project g_s into the same global illumination of the model texture. With a light field AAM the normalized mean texture is a light field, and thus cannot be directly applied to normalize g_s in Equation (11). Instead, we normalize both $g_m = F(G_m, \epsilon)$ and g_s to have zero mean and unit variance. We found this normalization scheme to work well in our experiments.

6 Experiments

We built a light field morphable model of the human head by capturing light fields of 50 subjects using a real-time light field camera array [23]. We collected 48 views (6 x 8) of each individual and manually segmented the head from each light field. Our head database consists of 37 males and 13 females of various races. Of these people, 7 are bearded and 17 are wearing glasses. The images in each view of the prototype light fields have resolution 320 x 240. Within each image, the head spans a region of approximately 80 x 120 pixels. The field of view captured by the camera array is approximately 25 degrees horizontally and 20 degrees vertically. To perform feature tracking, as described in Section 3, we used a multi-resolution Lukas-Kanade optical flow algorithm [15], with 4 pyramid levels and Laplacian smoothing¹. For comparison, we built a view-based AAM using the views of the light field camera array [7]. In both the definition of the view-based and light field active appearance models the parameter perturbations displayed in Table 1 were used to numerically compute the Jacobian matrix. To avoid over-fitting to noise, texture-shape PCA vectors having low variance were

¹ We acknowledge Tony Ezzat for the Lukas-Kanade optical flow implementation.

Table 1. Perturbation scheme used in both the view-based and light field AAMs. [20]

Variables	Perturbations
x, y	$\pm 5\%$ and $\pm 10\%$ of the height and width of the reference shape
θ	$\pm 5, \pm 15$ degrees
$scale$	$\pm 5\%, \pm 15\%$
c_{1-k}	$\pm 0.25, \pm 0.5$ standard deviations

discarded from each model, the remaining PCA vectors modelling 90% of the total model variance.

We implemented the view-based and light field active appearance models in MATLAB. To perform light field rendering we use the unstructured lumigraph algorithm described in [5]. In our experiments, our matching algorithm typically converged between 4 and 15 iterations when matching to an image and between 4 and 10 iterations when matching to a light field. Each iteration took a few seconds in un-optimized MATLAB. We believe that using a real-time light field renderer [5] would result in matching times similar to those reported for a 2D AAM [20].

7 Results

In this section we provide a comparison between a light field and a 2D view-based active appearance model. We then present various model matching experiments using our head light field appearance manifold.

Comparison to a View-Based AAM. To compare our method to a view-based AAM we built a single-view 2D AAM and compared it against a light field AAM. Each model was constructed using all fifty subjects, and was matched to a side view of two people. The resulting fits are displayed in Figure 3. In this figure one person is wearing glasses which self-occlude the subject in extreme views of the camera array. These self-occlusions are difficult to model using a view-based AAM, where inter-pose variation is modelled as shape. Also note that the view-dependent texturing effects in the persons glasses are preserved by the light field AAM, but are lost by the view-based AAM even though the person remains in the model.

Model Matching. To demonstrate the ability to fit a light field AAM to a single 2D image or light field, we match a novel person to the constructed head manifold using “leave-one-out” experimentation. Figure 4 illustrates fitting light fields of two people taken out of the model. To conserve space, only select views of each light field are displayed. Both fits are shown superimposed onto the corresponding input light field. Each light field is also provided for ground truth comparison. As seen from the figure, the input light fields are well matched and a convincing reconstruction of each person is generated. Specifically, the shape and texture of both individuals is well captured across views.

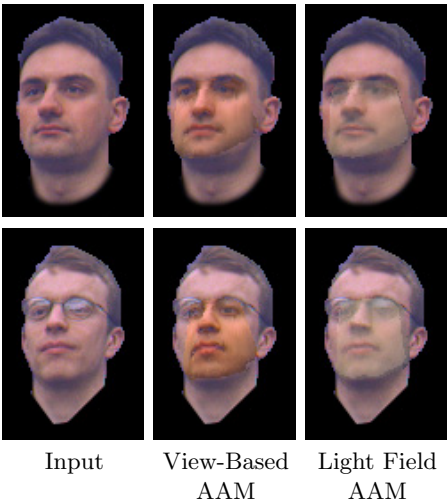


Fig. 3. Comparison of a light field active appearance model to a view-based AAM. The left column shows the input, the middle column the best fit with a 2D AAM, and the right column the light field fit. The 2D and light field appearance models both exhibit qualitatively good fits when the surface is approximately smooth and lambertian. When glasses are present, however, the 2D method fails and the light field appearance model succeeds.

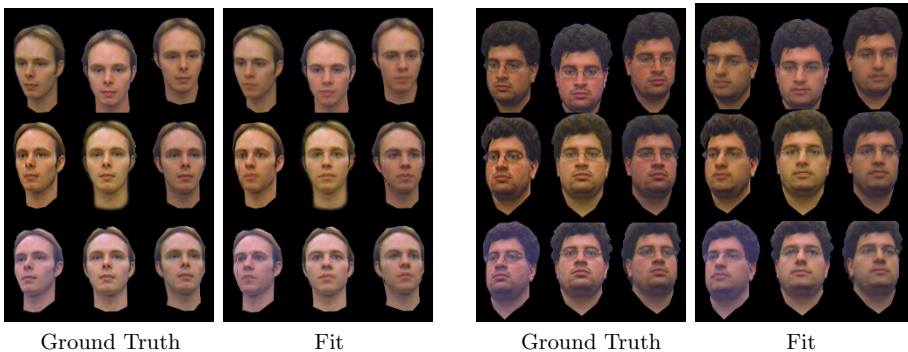


Fig. 4. Matching a light field AAM to a light field of a novel subject.

Figure 5 illustrates our model’s ability to generate convincing light field reconstructions from 2D images. This figure provides two example matches to 2D images with known pose. For each match, the person was removed from the model and imaged at a randomly selected pose not present in the light field AAM. The fit, rendered at the selected pose of each person, is displayed below each input image. The fitted light fields are also displayed. Note our method built a light field with 48 views from a single 2D image.

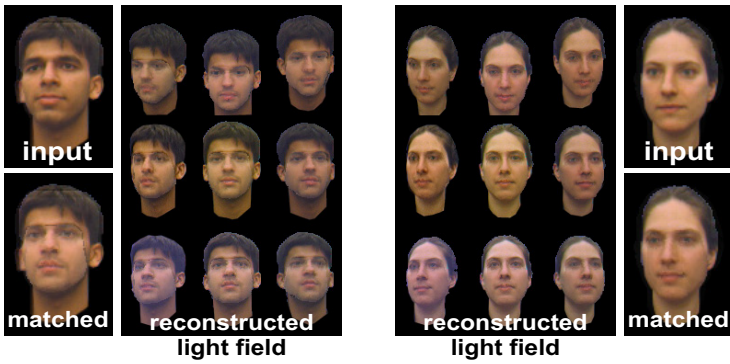


Fig. 5. Matching a light field AAM to 2D images of novel subjects. Each person is matched at a known pose. The reconstructed light field, is rendered over the input view and is displayed aside each match. The light field appearance model generates convincing light field reconstructions from 2D images. In particular, the overall shape and texture of each subject are well approximated across each view.

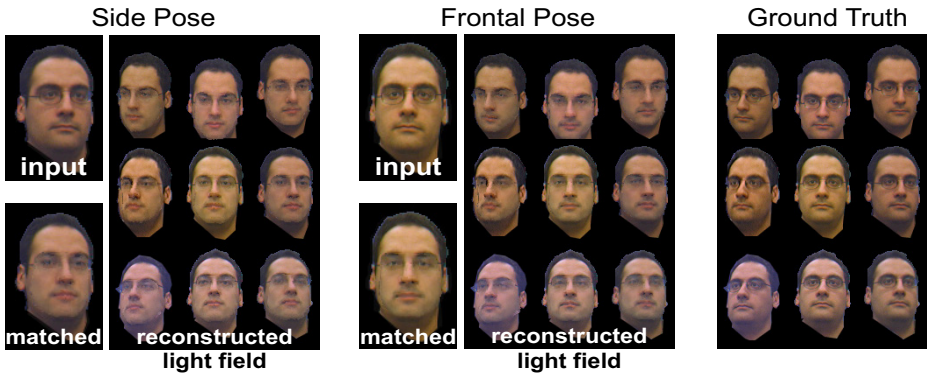


Fig. 6. Matching a light field AAM using automatic pose estimation (side pose). A match to a frontal, known pose is also provided for comparison. Note the reconstructed light fields are the same for both poses. Ground truth is shown on the right for comparison.

Figure 6 displays a fit to the head model using an unknown view of a person, in which pose was automatically estimated as described in Section 5. The model was also matched to a frontal view to verify that the reconstructed light fields are independent of input pose. As before this person is removed from the model and the views are not present in the light field AAM. The extreme views of the model light field fits are overlaid onto a captured light field of the subject. This light field is also shown as ground truth. Comparing each fit one finds that although the characteristics of the matched views are favored, the reconstructed light fields are strikingly similar. Also, note the view-dependent texturing effects present in

the subjects glasses, captured by the model. Comparing the matches of the above figure, one finds that our algorithm performs well in matching novel light fields and 2D images to the head manifold. Namely, the skin color, facial hair, and overall shape and expression of each novel subject are well approximated.

8 Conclusion and Future Work

We introduced a novel active appearance modeling method based on an image-based rendering technique. Light field active appearance models overcome many of the limitations presented by current 2D and 3D appearance models. They easily model complex scenes, non-lambertian surfaces, and view variation. We demonstrated the construction of a light field manifold of the human head using 50 subjects and showed how to match the model to a light field or single 2D image of a person outside of the model. In future work we hope to construct a camera array with a wider field of view that utilizes a non-planar camera configuration. We expect our approach to scale directly to the construction of dynamic light-field appearance manifolds, since our capture apparatus works in real-time.

References

1. E. H. Adelson and J. Bergen, *Computation Models of Visual Processing*. Cambridge: MIT Press, 1991, ch. The Plenoptic Function and the Elements of Early Vision.
2. D. Beymer and T. Poggio, "Face recognition from one example view, Tech. Rep. AIM-1536, September 1995.
3. V. Blanz and T. Vetter, "A morphable model for the synthesis of 3D faces," in *SIGGRAPH*, Los Angeles, 1999, pp. 187–194.
4. M. Brand, "Morphable 3D models from video," in *CVPR*, May 2001.
5. C. Buehler, M. Bosse, L. McMillan, S. J. Gortler, and M. F. Cohen, "Unstructured lumigraph rendering," in *SIGGRAPH*, 2001, pp. 425–432.
6. T. F. Cootes, G. J. Edwards, and C. J. Taylor, "Active appearance models," *Lecture Notes in Computer Science*, vol. 1407, pp. 484–98, 1998.
7. T. F. Cootes, G. V. Wheeler, K. N. Walker, and C. J. Taylor, "View-based active appearance models," *Image and Vision Computing*, vol. 20, pp. 657–664, 2002.
8. G. Edwards, C. Taylor, and T. Cootes, "Interpreting face images using active appearance models," in *3rd International Conference on Automatic Face and Gesture Recognition*, 1998, pp. 300–305.
9. P. Fua and C. Miccio, "From regular images to animated heads: a least squares approach," in *ECCV*, Springer, Berlin, 1999, pp. 188–202.
10. C. Goodall, "Procrustes methods in the statistical analysis of shape." *Journal of the Royal Statistical Society*, vol. 53, no. 2, pp. 285–339, 1991.
11. S. J. Gortler, R. Grzeszczuk, R. Szeliski, and M. F. Cohen, "The lumigraph," *Computer Graphics*, vol. 30, no. Annual Conference Series, pp. 43–54, 1996.
12. R. Gross, I. Matthews, and S. Baker, "Appearance-based face recognition and light fields," *IEEE PAMI*, vol. 26, no. 4, April 2004.

13. M. J. Jones and T. Poggio, "Multidimensional morphable models," in *ICCV*, 1998, pp. 683–688.
14. M. Levoy and P. Hanrahan, "Light field rendering," *Computer Graphics*, vol. 30, pp. 31–42, 1996.
15. B. D. Lucas and T. Kanade, "An iterative image registration technique with an application to stereo vision," in *International Joint Conference on Artificial Intelligence*, 1981, pp. 674–679.
16. B. Moghaddam and A. Pentland, "Probabilistic visual learning for object recognition," *IEEE PAMI*, vol. 19, no. 7, pp. 696–710, 1997.
17. H. Murase and S. Nayar, "Visual learning and recognition of 3-d objects from appearance," *IJCV*, vol. 14, no. 1, pp. 5–24, 1995.
18. F. H. Pighin, R. Szeliski, and D. Salesin, "Resynthesizing facial animation through 3d model-based tracking," in *ICCV*, 1999, pp. 143–150.
19. S. Romdhani, S. Gong, and A. Psarrou, "A multi-view nonlinear active shape model using kernel pca," in *British Machine Vision Conference*, 1999, pp. 483–492.
20. M. B. Stegmann, "Analysis and segmentation of face images using point annotations and linear subspace techniques," Technical University of Denmark, DTU, Tech. Rep., August 2002.
21. M. Turk and A. Pentland, "Eigen faces for recognition," *Journal of Cognitive Neuroscience*, vol. 3, no. 1, 1991.
22. B. Wilburn, M. Smulski, H.-H. K. Lee, and M. Horowitz, "The light field video camera," in *Proceedings of Media Processors 2002, SPIE Electronic Imaging*, 2002.
23. J. C. Yang, M. Everett, C. Buehler, and L. McMillan, "A real-time distributed light field camera," in *Eurographics Workshop on Rendering*, 2002, pp. 1–10.
24. Z. Zhang, L. Wang, B. Guo, and H.-Y. Shum, "Feature-based light field morphing," in *Conference on Computer graphics and interactive techniques*. ACM Press, 2002, pp. 457–464.