



HOKKAIDO UNIVERSITY

Title	A Component-based Face Synthesizing Method
Author(s)	Chiang, Cheng-Chin; Chen, Zhih-Wei; Yang, Chen-Ning
Description	APSIPA ASC 2009: Asia-Pacific Signal and Information Processing Association, 2009 Annual Summit and Conference. 4-7 October 2009. Sapporo, Japan. Oral session: 3D Synthesis and Expression (5 October 2009).
Citation	Proceedings : APSIPA ASC 2009 : Asia-Pacific Signal and Information Processing Association, 2009 Annual Summit and Conference, 24-30
Issue Date	2009-10-04
Doc URL	https://hdl.handle.net/2115/39638
Type	conference paper
File Information	MA-L1-4.pdf



A Component-based Face Synthesizing Method

Cheng-Chin Chiang, Zhih-Wei Chen, and Chen-Ning Yang

National Dong Hwa University,

Dept. of Comput. Sci. and Info. Eng. Shoufeng, Hualien, Taiwan

E-mail: ccchiang@mail.ndhu.edu.tw Tel: +886-3-8634027

Abstract—The active appearance models (AAM) is a popular tool in object tracking. An AAM is featured by its integrated modeling of deformations in both shapes and textures. Therefore, in addition to the object tracking, AAM is also a good visual synthesizer. The other strength of the AAM is its compact representations for the geometries and textures of synthesized objects. By training with the principal component analysis method, the AAM parameterizes the shape and the texture of each synthesized object simply with a linear combination of eigen-shapes and eigen-textures, respectively. This paper presents a novel video-driven face synthesizing method which tracks the faces of a person on video frames and synthesize novel faces using geometries of individual facial components, such as eyes, noses, and mouths, of other persons. To this end, we propose the component-based active shape models (ASM) for synthesizing each facial component. One major prominent feature of the proposed method is that a rich number of novel facial expressions can be synthesized by combining different facial components from different persons on the synthesized faces. No further retraining process for the AAMs or ASMs is required for synthesizing these novel facial expressions. The experimental results show that the proposed method can accomplish interesting and vivid facial synthesis and exhibits its high potential in many practical applications.

I. INTRODUCTION

Computer facial animation is an interesting problem in computer graphics and has high potential in many applications including movie producing, computer games, computer-aided instruction, and so on. From the perspective of human communication, a communication interfaced by computer facial animation attains much higher attractiveness than does a communication interfaced by text-only outputs [10]. Driven by the attractiveness, many researchers devoted to this research. Particularly, real-time facial animation driven by faces on video sequences lures extensive attention from researchers. Many different works have been proposed for facial animation [1] [15] [13]. Among these many works, a general approach is the feature-based facial animation which demands the definition and then the positioning of salient feature points on faces. Therefore, the performance strongly relies on the characteristics of feature points including the number and the accuracy of positioned feature points. Unfortunately, locating features points on faces is not trivial because many situations, such as varying illumination and self-occlusion, further complicate this problem. In fact, handling these situations for robustly locating facial feature points is still a big challenge.

Lin [9] proposed a 3D model-based approach using multi-view face images. He simplified the task to locating feature points by putting artificial markers on faces. The geometry relationships between the multi-views are known in advance and the environmental conditions are well controlled. The positions of each marker on different views are located to acquire the 3D trajectory. The Kalman filter is also applied to robustly track the markers. Dutreue [6] manually set the initial locations of feature points and used pyramidal Lucas-Kanade tracker to track the selected points. A generic 3D face model is then warped via RBFs (radial basis functions) using the tracked feature points as control points. The motion of the head and facial components are assumed with limited in-plane rotations to attain good tracking results. In contrast to the 3D approach, Chuang [2] uses the "Eigen-Points" [5] algorithm to track the 2D motion of facial features. The tracked facial features on each frame define a key-shape for the face. The facial animations for in-between frames are synthesized by interpolation with key-shape blending.

The active appearance model (AAM) [3] is a popular tool for object tracking. More and more works [7], [14], [12] use AAMs to track faces and synthesizes facial expressions with parameters estimated from the AAM tracking. For example, Jiang [8] used the AAM to track feature points defined in MPEG-4 FDPs on frontal face. The tracked FDPs then drive the facial animation. This approach is suitable for facial animation over low-bandwidth networks.

Theobald [12] animated facial expressions by using the AAM. The synthesis is accomplished by two AAMs, one for tracking the source face and the other for synthesizing the target face. The source face and target face may be from different persons. After tracking the source face, applying the shape displacements of the tracking AAM onto the synthesizing AAM obtains the synthesized facial expression.

Our approach is similar to Theobald's method with one major difference. In our approach, we propose the component-based active shape model (ASM) [4] to replace the whole-face AAMs. The proposed component-based synthesizer benefits that novel faces can be easily synthesized by combining facial components from different persons without any further training. However, Theobald's work is limited to synthesize trained faces only. New training processes are required if novel facial expressions are to be synthesized.

The rest of this paper is organized as follows. In Section 2, we briefly describe the concept of AAM. The component-based facial animation synthesizer is presented in Section 3.

Section 4 demonstrates some results of synthesized facial expressions. Finally, Section 5 ends this paper with some concluding remarks.

II. ACTIVE APPEARANCE MODEL (AAM)

The active appearance model (AAM) [3] is a popular tool for object tracking. An AAM is featured by its integrated modeling of deformations in both shapes and textures of objects. Therefore, in addition to the object tracking, the AAM is also a good synthesizer for visual objects. The other strength of the AAM is its compact representations for the geometries and textures of synthesized objects. Using the principal component analysis (PCA), the AAM parameterizes the shape and the texture of each synthesized object with a linear combination of eigen-shapes and eigen-textures, respectively. Given sufficient training samples, an AAM can well handle deformations in both shapes and textures of objects to a large extent.

Basically, each object shape can be represented with a set of 2D landmarks whose X-Y coordinates are stored into a shape vector $\mathbf{s} = [x_1, y_1, \dots, x_i, y_i, \dots]$, where (x_i, y_i) is the coordinate of the i^{th} landmark. An object texture a can thus be obtained by performing a warping function $W()$ on the object image $\mathbf{I}$ with respect to a warping parameter vector $\mathbf{p}$, i.e., $a = \mathbf{I}(W(\mathbf{s}; \mathbf{p}))$. Therefore, different instances of the warping parameter vector $\mathbf{p}$ would lead to different appearances, including both the shape and the texture, of the object.

The basic principle underlying the AAM is that the shape and the texture of an object can be represented with linearly-combined basis shapes and basis textures, respectively. These basis shapes and the basis textures are called eigen-shapes and eigen-textures, respectively. Therefore, before using an AAM, the first task is to find these eigen-shapes and eigen-textures through the PCA. Briefly speaking, given a set of shape vectors $T_S = \{\mathbf{s}_1, \mathbf{s}_2, \dots, \mathbf{s}_N\}$, a covariance matrix can be derived by

$$\Sigma_S = \frac{1}{N-1} \sum_{k=1}^N (\mathbf{s}_k - \mathbf{S}_0)(\mathbf{s}_k - \mathbf{S}_0)^t, \quad (1)$$

where $\mathbf{S}_0 = \frac{1}{N} \sum_{k=1}^N \mathbf{s}_k$. Consequently, the eigen-shapes can be obtained by finding the eigenvectors of Σ_S that correspond to largest M eigenvalues, where M is the number of used eigen-shapes and is specified by the designer. In general, M is set to a small value for reducing the amount of required data storages and computations. To find the eigen-textures, each collected texture samples is first warped to the mean shape defined by $\mathbf{S}_0$. Afterwards, the eigen-textures can be obtained in a way similar to that of obtaining eigen-shapes using the set of warped texture samples.

With the obtained eigen-shapes $S = [\mathbf{S}_1, \dots, \mathbf{S}_M]$ and eigen-textures $A = [\mathbf{A}_1, \dots, \mathbf{A}_{M'}]$, an object shape and its texture can be represented by the linear combination of the eigen-shapes and eigen-textures as follows:

$$a = \mathbf{A}_0 + \lambda_1 \mathbf{A}_1 + \dots + \lambda_{M'} \mathbf{A}_{M'}, \quad (2)$$

$$\mathbf{s} = \mathbf{S}_0 + \beta_1 \mathbf{S}_1 + \dots + \beta_M \mathbf{S}_M, \quad (3)$$

where $\mathbf{A}_0$ and $\mathbf{S}_0$ are the mean texture and the mean shape, respectively.

When tracking the object on an image frame with an AAM, the goal is to estimate the warping parameter vector $\mathbf{p}$ so that the linearly-combined textures best fits the image obtained by warping the current object image to the mean shape with respect to $\mathbf{p}$. That is, the following error must be minimized:

$$\sum_{\mathbf{x} \in \mathbf{S}_0} [\mathbf{A}_0(\mathbf{x}) + \sum_{i=1}^{M'} \lambda_i \mathbf{A}_i(\mathbf{x}) - \mathbf{I}(W(\mathbf{x}, \mathbf{p}))]^2. \quad (4)$$

Usually, an appearance warping can be modelled with a local warping $W()$ and a global warping $N()$. Corresponding to these two kinds of warping, two warping parameter vectors, denoted respectively as $\mathbf{p}$ and $\mathbf{q}$, are used. Matthews [11] thus proposed an accelerated version of AAM called inverse compositional AAM (ICAAM) which modified the error measure as

$$\sum_{\mathbf{x} \in \mathbf{S}_0} [\mathbf{A}_0(\mathbf{x}) + \sum_{i=1}^{M'} \lambda_i \mathbf{A}_i(\mathbf{x}) - \mathbf{I}(N(W(\mathbf{x}, \mathbf{p}); \mathbf{q}))]^2. \quad (5)$$

Matthews assumed the global warping $N()$ to the similarity transform in their implementation. During the AAM fitting process, the AAM incrementally updates the warping parameter vectors until the error defined in Eq. (5) ceases decreasing. After obtaining parameter vectors $\mathbf{p}$ and $\mathbf{q}$, the coefficients λ_i are computed by

$$\lambda_i = \sum_{\mathbf{x} \in \mathbf{S}_0} \mathbf{A}_i(\mathbf{x}) \cdot [\mathbf{I}(N(W(\mathbf{x}, \mathbf{p}); \mathbf{q})) - \mathbf{A}_0(\mathbf{x})]. \quad (6)$$

Finally, the estimated shape warping parameter vectors, $\mathbf{p}$ and $\mathbf{q}$, and the texture coefficients λ_i can then be used to synthesize the appearance of the tracked object.

III. FACIAL SYNTHESIS BY COMPONENT-BASED ACTIVE SHAPE MODEL

The AAM has its tempting features in performing the task of object tracking. Thanks to the compact representation through the linear combination of eigen-shapes and eigen-textures, the AAM is also a good choice for synthesizing facial animation. However, the major weakness of the AAM is that the whole synthesized face is tightly coupled with the face appearance of some specific person. The synthesis would be not possible to mix up different local facial components from different persons.

The approach proposed in this paper exploits a novel method for enabling the mixing of different facial components from different persons when synthesizing facial animation. We propose the component-based active shape models (ASM) for synthesizing each facial component independently. This way is beneficial to synthesize novel faces composed by facial components of different persons with no need to retrain the AAM. Hence, due to the many combinations of facial components from different persons, a rich number of novel faces can be synthesized by the proposed method. In the following, we present the details of the component-based ASM and describe how to synthesize facial animation.

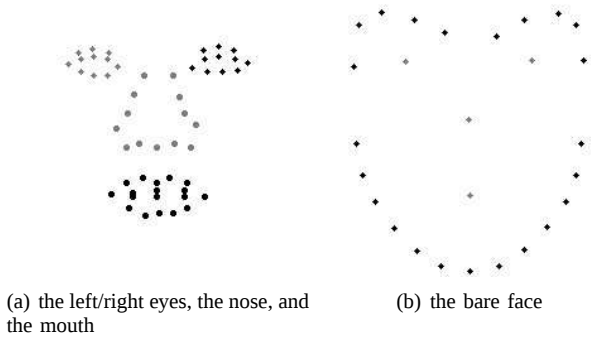


Fig. 1. The vertices of the facial components.

A. Component-based ASM training

A component-based ASM is actually an AAM with only the mean shape and eigen-shapes. For each facial component of every person, we build an ASM. These built ASMs provide the geometric information for aligning the shapes of the corresponding facial components from different persons. Therefore, the facial expressions of a person A can be creatively synthesized by aligning the shapes of some facial components to those of the corresponding facial components of other persons.

The training of a component-based ASM is similar to that of an AAM. Let $F_\Omega = \{x_1, y_1, \dots, x_n, y_n\}$ denote the vertices that represent the facial landmarks. As shown in Fig. 1(a) and Fig. 1(b), we divide the face into five components, include the left eye, the right eye, the nose, the mouth, and the bare face with each facial component identified by a landmark at the component center. For convenience, we denote the five disjoint sets of landmarks as $F_{\phi \in \{leye, reye, nose, mouth, bare\}}$. Notice that the eye components, F_{leye} and F_{reye} , contain the eyelid. During the training process, the training samples are firstly grouped, by persons, into several different sets of component landmarks, denoted with $\mathbf{F}_\phi^i$, where i and ϕ are respectively the person identity and the component identity. By applying the PCA on each $\mathbf{F}_\phi^i$, we derive the mean-shape and the eigen-shape vectors $\mathbf{S}_\phi^i = \{\mathbf{S}_{\phi,0}^i, \mathbf{S}_{\phi,1}^i, \dots, \mathbf{S}_{\phi,m_\phi,i}^i\}$, where $m_{\phi,i}$ is the number of used eigen-shapes which may vary with different i and ϕ . Therefore, a novel mixing-up face can be defined with $\mathbf{S}^{new} = \{\mathbf{S}_{leye}^i, \mathbf{S}_{reye}^j, \mathbf{S}_{nose}^k, \mathbf{S}_{mouth}^l, \mathbf{S}_{base}^m\}$, for different combinations of i, j, k, l , and m .

B. Component-based facial expression syntheses

Six steps are required in our facial animation synthesis. The first step is to track the face of the person with the identity tr on the input video frame using the AAM. The second step is to separately extract each facial component from the tracked face. On each extracted facial component, the displacement from the tracked position of each landmark to that of the corresponding landmark on the mean shape is computed, denoted as ΔF_ϕ^{tr} . In the third step, we perform the component-based shape alignment by projecting the computed displacements onto the corresponding component eigenspace

of the source person with another identity new , i.e.,

$$\Delta P_\phi^{new} = (\Phi_\phi^{new})^T \cdot \Delta F_\phi^{tr} = [\Delta p_1, \dots, \Delta p_m]^T, \quad (7)$$

where $\Phi_\phi^{new} = [S_{\phi,1}^{new}, \dots, S_{\phi,m}^{new}]$ is matrix composed by the eigen-shapes of the source person's facial component. The fourth step is to infer the positions of the landmarks for the synthesized facial component by

$$F_\phi^{new} = S_{\phi,0}^{new} + \sum_{i=1}^m \Delta p_i \cdot S_{\phi,i}^{new}. \quad (8)$$

After computing the F_ϕ^{new} , the displacements of the five component centers in the bare face must be added to the corresponding F_ϕ^{new} . Eq. (8) finishes the local warping of facial components. As the face may be undergoing a global head motion which can also be estimated by the ICAAM, the fifth step is thus to perform the global transformation on the obtained F_ϕ^{new} according to the ICAAM's estimation. Finally, the last step is to synthesize the textures for the face. The texture of the synthesized face is obtained by directly warping the texture on the tracked face of the person tr .

IV. EXPERIMENTAL RESULTS

A. Data collection and model building

Several experiments are conducted to examine the effectiveness of the proposed method. We collect face images from six subjects for training the AAMs and ASMs. The number of collected image samples for each subject ranges from 13 to 17. Fig. 2 shows some samples of these six subjects, including four real humans and two artificial movie actors (Shrek and Gollum). The face samples are captured from difference sources, including photos captured from low-cost web cameras and movie clips. For the artificial movie actors, the collected images may lack samples for some facial expressions. This lack of facial expressions could affect the tracking and the syntheses of some expressions to some extent.

For each subject, an AAM is trained for person-specific face tracking. Simultaneously, an ASM is also trained for each facial components of every subject. Fig. 3(a) and Fig. 3(b) show the mean shape and the mean texture of each subject used in the corresponding AAM, respectively.

B. Synthesized results

From Fig. 4 to Fig. 9, we demonstrate some synthesized results for novel faces. These novel faces are composed by assembling facial components of different geometries from different subjects. In each figure, the sub-figure (a) shows the original image and the sub-figure (b) shows the tracked face shape by the person-specific AAM. Sub-figures (c) and (d) show the synthesized face image and the corresponding shape, respectively. Note that we use the bare face of the tracked subject as the synthesized bare face in these experiments. Therefore, the synthesized face can be seamlessly combined with the background including the hairs. If a component ASM can also be built for the hairs, then the synthesized faces would be more vivid. In these demonstrated results,

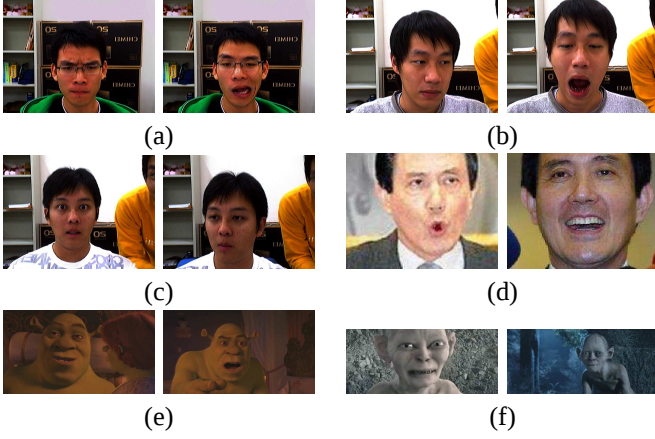


Fig. 2. Some samples of six subjects: (a) Subject "Ning"; (b) Subject "Edward"; (c) Subject "Joker"; (d) Subject "Ma"; (e) Subject "Shrek"; (f) Subject "Gollum".

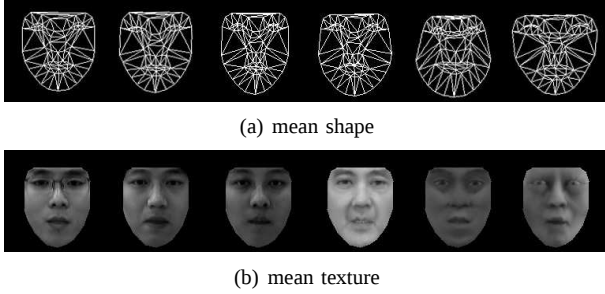


Fig. 3. The mean shape and the mean texture for each subject.

each synthesized novel face combines the facial components from three to four different subjects. The results reveal that a rich number of novel facial expressions can be synthesized using the proposed component-based synthesizing method. For example, the synthesized face in Fig. 4 having a smaller nose and a smaller mouth evidently differs from the original face of Subject Ning. In Fig. 7, Ma becomes older due to the synthesized wrinkled facial texture from Gollum. Since Shrek's nose is large, the synthesized result in Fig. 5 leads to an interesting face with a big nose. In Fig. 6, Edward's mouth is enlarged by using the geometry of Gollum's mouth. Shrek's face in Fig. 8 becomes more human-like after the synthesis. The facial expression of Gollum in Fig. 9 changes from an angry face to a sad one by reducing both the sizes of eyes and mouth. From Fig. 10 and Fig. 13, we demonstrate some synthesized results of using one subject's whole facial components.

V. CONCLUDING REMARKS

This paper proposes a novel synthesizing method to composing facial expressions from video inputs. Different from the feature-based approach which requires the positioning of salient features on faces, we use a person-specific AAM to track the whole face of the user. With the AAM, the face tracking problem turns into an iterative error minimization process to estimate the best warping parameters of face shapes.

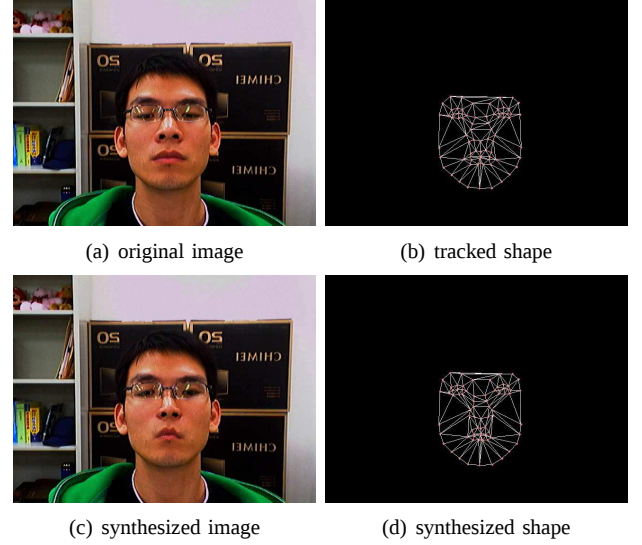


Fig. 4. A novel face composed with Leye (from Shrek), Reye (from Shrek), Nose (from Gollum), Mouth (from Joker) and Bare Face (from Ning).

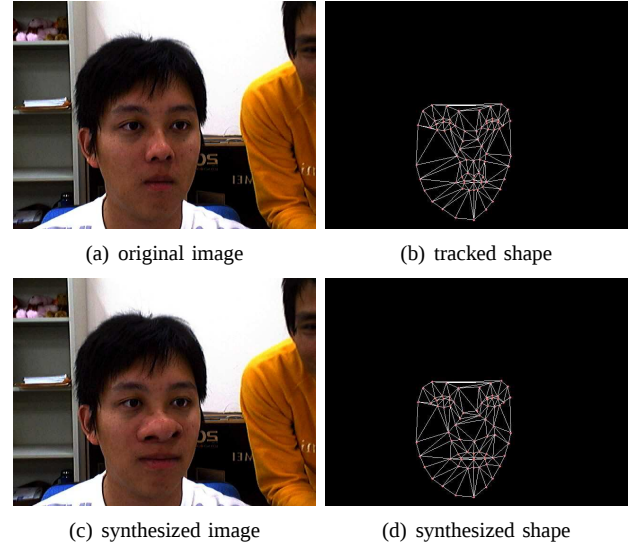


Fig. 5. A novel face composed with Leye (from Ma), Reye (from Ma), Nose (from Shrek), Mouth (from Edward) and Bare Face (from Joker).

From the estimated warping parameters, the warping of facial components in the AAM can be isolated and aligned with the geometries of facial components of other persons. Therefore, the proposed method can synthesize a rich number of novel faces with no need to retrain any AAMs or ASMs. This capability is not possible for other synthesizing methods that perform whole-face syntheses of facial expressions. From the experimental results, the proposed method accomplishes interesting and vivid syntheses of various facial expressions. The results demonstrate the high potential of the proposed method in many practical applications.

In the future, the proposed work may be extended towards the following directions:

- the extension of the person-specific AAMs to person-

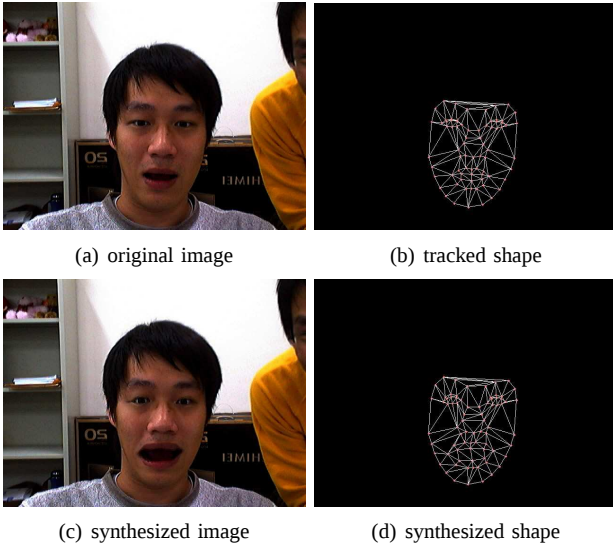


Fig. 6. A novel face composed with Leye (from Sherk), Reye (from Sherk), Nose (from Ning), Mouth (from Gollum), and Bare Face (from Edward).

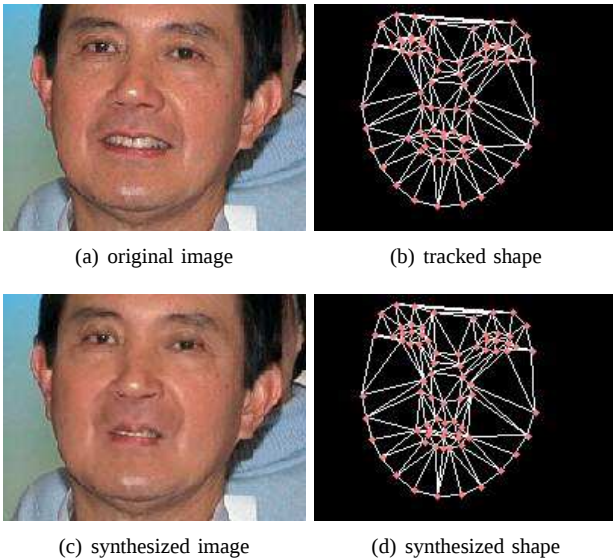


Fig. 7. A novel face composed with Leye (from Ning), Reye (from Ning), Nose (from Gollum), Mouth (from Joker), and Bare Face (from Ma).

independent AAMs

- the inclusion of other facial components such as hairs, teeth, and ears;
- the development of 3D component-based ASM models;

This work creates a new shape for a new character, but the texture directly uses the tracked face. In the future, we will try to create a new texture by associating different person's face texture. This isn't easy work, because the face textures of difference characters are varied. The synthesis work of wrinkle and dimple are the new challenges in the new texture synthesis mechanism. Finally, to create a realistic 3d facial animation in our future work.

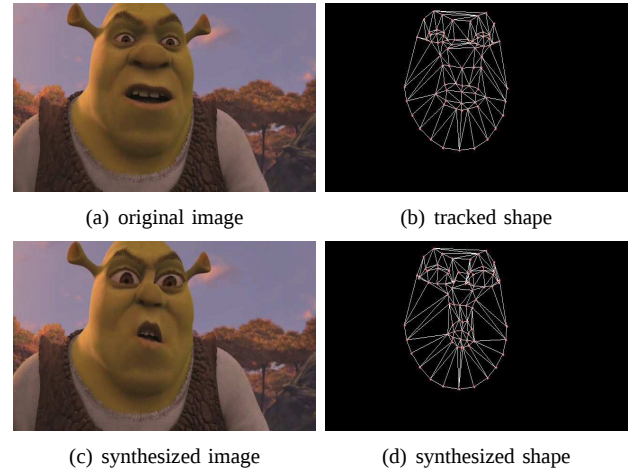


Fig. 8. A novel face composed with Leye (from Gollum), Reye (from Gollum), Nose (from Ma), Mouth (from Joker), and Bare Face (from Shrek).

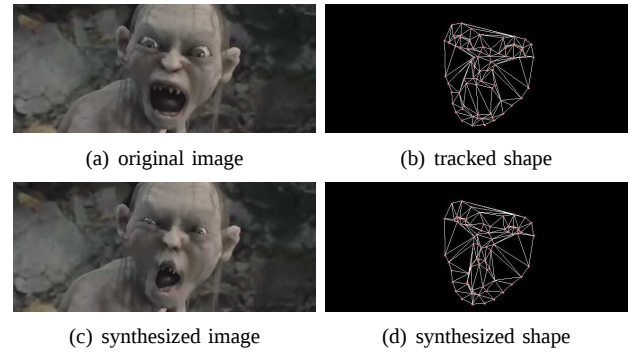


Fig. 9. A novel face composed with Leye (from Edward), Reye (from Edward), Nose (from Joker), Mouth (from Joker), and Bare Face (from Gollum).

REFERENCES

- [1] B. Choe, H. Lee, and H.S. Ko. Performance-driven muscle-based facial animation. *The Journal of Visualization and Computer Animation*, 12(2):67–79, 2001.
- [2] E. Chuang and C. Bregler. Performance driven facial animation using blendshape interpolation. *Computer Science Technical Report, Stanford University*, pages 2002–02, 2002.
- [3] T.F. Cootes, G.J. Edwards, C.J. Taylor, et al. Active appearance models. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 23(6):681–685, 2001.
- [4] T.F. Cootes, C.J. Taylor, D.H. Cooper, J. Graham, et al. Active shape models-their training and application. *Computer vision and image understanding*, 61(1):38–59, 1995.
- [5] M. Covell and C. Bregler. Eigen-points. In *Proc. Int. Conf. Image Processing*, pages 471–474, 1996.
- [6] L. Dutreue, A. Meyer, and S. Bouakaz. Feature points based facial animation retargeting. In *Proceedings of the 2008 ACM symposium on Virtual reality software and technology*, pages 197–200. ACM New York, NY, USA, 2008.
- [7] G. Fanelli and M. Fratarcangeli. A Non-Invasive Approach for Driving Virtual Talking Heads from Real Facial Movements. In *3DTV Conference, 2007*, pages 1–4, 2007.
- [8] B. Jiang, J. Bu, and C. Chen. Improving Visual Awareness by Real-Time 2D Facial Animation for Ubiquitous Collaboration. In *Computer Supported Cooperative Work in Design, 2007. CSCWD 2007. 11th International Conference on*, pages 151–156, 2007.
- [9] I.C. Lin, J.S. Yeh, and M. Ouhyoung. Extracting 3d facial animation parameters from multiview video clips. *IEEE Computer Graphics and Applications*, 22(6):72–80, 2002.

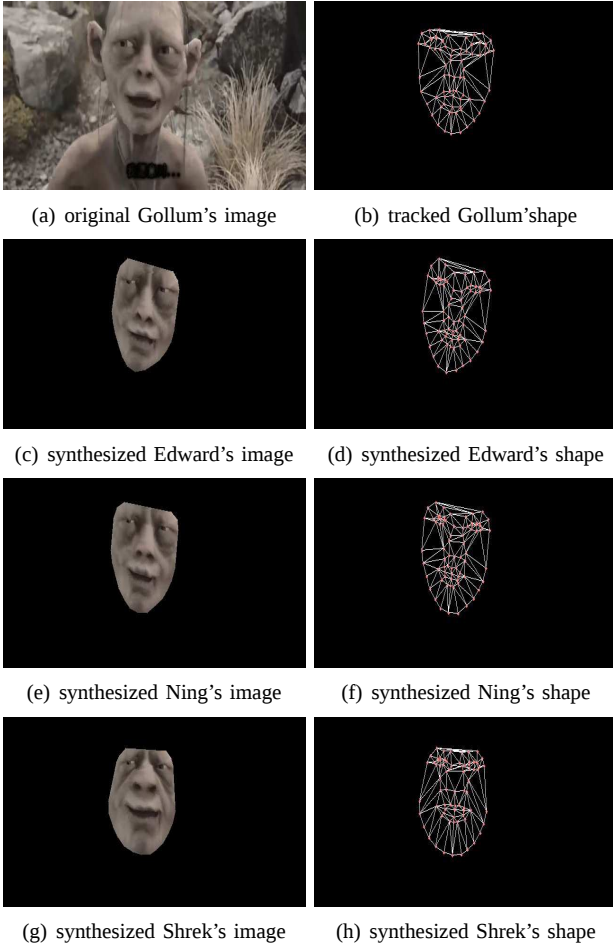


Fig. 10. Synthesis gollum's face by using different subject's whole components.

- [10] K. Liu, A. Weissenfeld, J. Ostermann, and X. Luo. Robust AAM building for morphing in an image-based facial animation system. In *2008 IEEE International Conference on Multimedia and Expo*, pages 933–936, 2008.
- [11] I. Matthews and S. Baker. Active appearance models revisited. *International Journal of Computer Vision*, 60(2):135–164, 2004.
- [12] B. Theobald, I. Matthews, J. Cohn, and S. Boker. Real-time expression cloning using active appearance models. In *Proceedings of the ACM International Conference on Multimodal Interfaces (ICMI'07)*, pages 134–139, 2007.
- [13] T. Vetter and V. Blanz. A morphable model for the synthesis of 3D faces. In *Proceedings of the ACM SIGGRAPH Conference on Computer Graphics*, pages 187–194, 1999.
- [14] SN Yang, KS Huang, YC Tseng, and CY Wang. AVATAR FACIAL ANIMATION SYNTHESIZER FOR MOBILE CLIENTS. In *Advances in Mobile Multimedia, 2004. MoMM 2004, 2th International Conference on*, 2004.
- [15] Q. Zhang, Z. Liu, G. Quo, D. Terzopoulos, and H.Y. Shum. Geometry-driven photorealistic facial expression synthesis. *IEEE Transactions on Visualization and Computer Graphics*, 12(1):48–60, 2006.

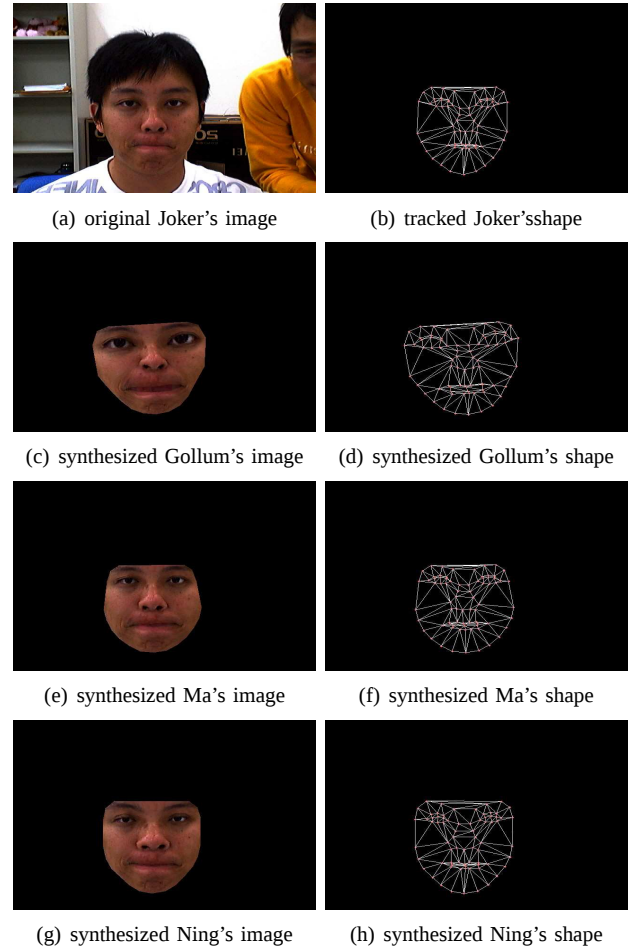
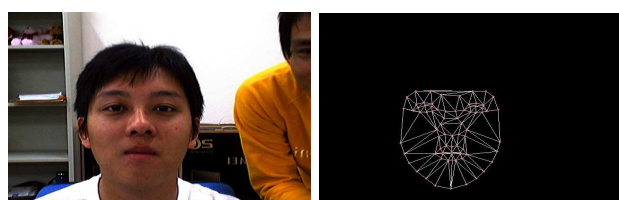
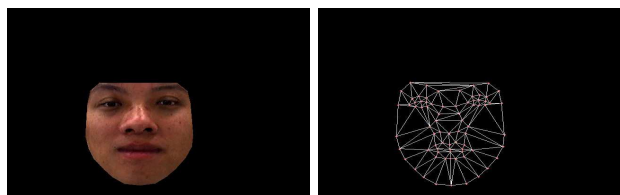


Fig. 11. Synthesis Joker's face by using different subject's whole components.



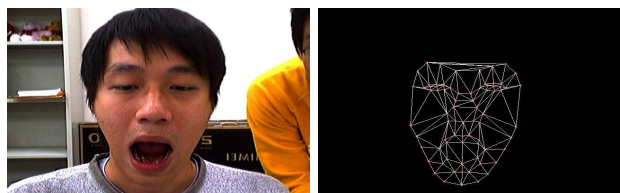
(a) original Joker's image

(b) Joker's shape



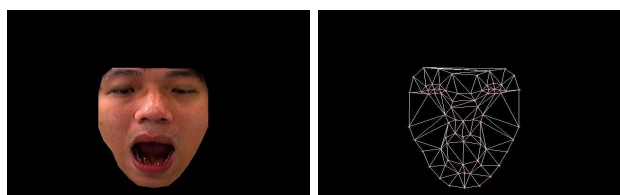
(c) synthesized img

(d) synthesized shape



(e) original Edward's image

(f) Edward's shape



(g) synthesized image

(h) synthesized shape



(i) original Shrek's image

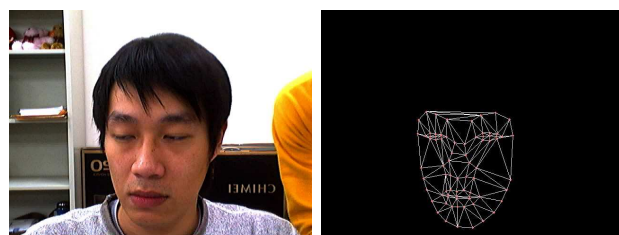
(j) Shrek's shape



(k) synthesized image

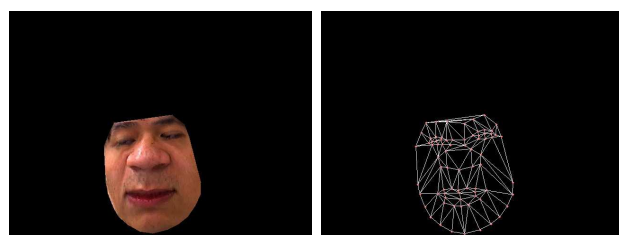
(l) synthesized shape

Fig. 12. Face synthesis using Ning's whole components



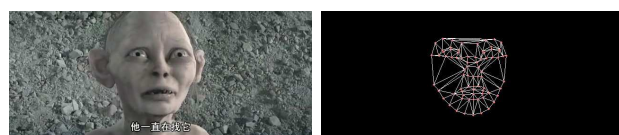
(a) original Edward's image

(b) Edward's shape



(c) synthesized image

(d) synthesized shape



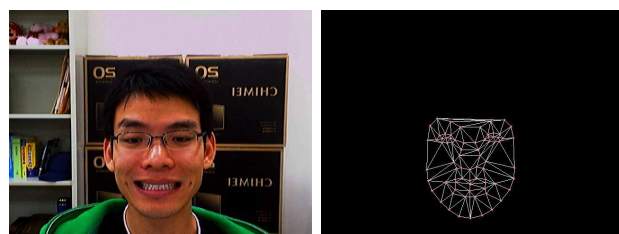
(e) original Gollum's image

(f) Gollum's shape



(g) synthesized image

(h) synthesized shape



(i) Ning's image

(j) Ning's shape



(k) synthesized image

(l) synthesized shape

Fig. 13. Face synthesis using Shrek's whole components.