

Cross-Pose Face Recognition — A Virtual View Generation Approach Using Clustering Based LVTM

Xi LI^{†a)}, *Nonmember*, Tomokazu TAKAHASHI^{†,††}, Daisuke DEGUCHI^{†††}, *Members*,
Ichiro IDE[†], *Senior Member*, and Hiroshi MURASE[†], *Fellow*

SUMMARY This paper presents an approach for cross-pose face recognition by virtual view generation using an appearance clustering based local view transition model. Previously, the traditional global pattern based view transition model (VTM) method was extended to its local version called LVTM, which learns the linear transformation of pixel values between frontal and non-frontal image pairs from training images using partial image in a small region for each location, instead of transforming the entire image pattern. In this paper, we show that the accuracy of the appearance transition model and the recognition rate can be further improved by better exploiting the inherent linear relationship between frontal-nonfrontal face image patch pairs. This is achieved based on the observation that variations in appearance caused by pose are closely related to the corresponding 3D structure and intuitively frontal-nonfrontal patch pairs from more similar local 3D face structures should have a stronger linear relationship. Thus for each specific location, instead of learning a common transformation as in the LVTM, the corresponding local patches are first clustered based on an appearance similarity distance metric and then the transition models are learned separately for each cluster. In the testing stage, each local patch for the input non-frontal probe image is transformed using the learned local view transition model corresponding to the most visually similar cluster. The experimental results on a real-world face dataset demonstrated the superiority of the proposed method in terms of recognition rate.

key words: face recognition, pose invariant, clustering, local view transition model

1. Introduction

Due to its wide range of potential real-life applications such as identity authentication, intelligent surveillance, human-computer interface and so on, face recognition has been one of the most active research topics in the biometric field within the computer vision and the pattern recognition communities [1]. Unlike other biometric techniques such as fingerprint recognition, palm print recognition or iris recognition, face recognition is inherently a passive and non-intrusive technique that has the advantage of not requiring cooperative subjects. That is to say, a practical face recognition system is supposed to have the ability to recognize the face of an uncooperative subject in an arbitrary situation and uncontrolled environment setting, even without the notice of the target subject. This advantage of environment setting

generality also poses great challenges to the problem of face recognition because as the viewing condition changes, the captured face appearances might vary too drastically to be easily identified. Within the past several decades, many face recognition methods have been proposed. However, most of those traditional methods can successfully recognize faces only when face images are captured under a constrained condition and a controlled environment, for example recognize frontal faces with normal expressions and typical indoor illuminations, which are usually unrealistic in many real-life application scenarios. Usually the performance of these traditional methods will degrade greatly when face images are captured in unconstrained conditions caused by factors such as varying viewpoints, illumination changes, occlusions, aging, expressions and poses.

This paper studies the problem of face recognition across poses, where each subject has a frontal gallery face image stored in the database and the probe image is not necessarily frontal. It is of great interest in many real-world face recognition application scenarios such as surveillance systems, where the subjects are either indifferent or uncooperative, so the captured face images are usually in low-resolution and non-frontal. Pose variation has been identified as one of the prominent difficult problems in the research of face recognition [1]. The major difficulty of the cross-pose face recognition is that the intra-person appearance differences caused by rotation are often larger than the inter-person differences. That is to say, the distance between appearance vectors of two faces of different persons under similar viewpoints is much smaller than that of the same person under different viewpoints. This phenomenon makes the traditional face recognition methods such as eigen-face [2] or Fisher-face [3] infeasible. Obviously, one straightforward method for cross-pose face recognition is to actively compensate pose variations by providing gallery views in each rotation angle to recognize rotated non-frontal probe views. This can be achieved by first collecting and preparing multiple real-view templates beforehand for every known individual in each specific pose condition. Although the number of required real gallery images can be reduced by proper quantization on the rotation angles due to the fact that general face recognition algorithms are able to tolerate small pose variations to some extent, the tedious process of collecting multiple face images in different poses for real-view based matching is still unfavorable and even impractical in some cases. For example, in the application of airport secu-

Manuscript received June 9, 2012.

Manuscript revised October 20, 2012.

[†]The authors are with the Graduate School of Information Science, Nagoya University, Nagoya-shi, 464–8601 Japan.

^{††}The author is with the Faculty of Economics and Information, Gifu Shotoku Gakuen University, Gifu-shi, 500–8288 Japan.

^{†††}The author is with Information and Communications Headquarters, Nagoya University, Nagoya-shi, 464–8601 Japan.

a) E-mail: murase@is.nagoya-u.ac.jp

DOI: 10.1587/transinf.E96.D.531

city surveillance systems, there is only one frontal passport photo per person that could be collected and stored in the database.

Previously both 3D model based methods [4], [5] and 2D appearance based methods [6]–[10] have been proposed for pose invariant face recognition. 3D Morphable Model [4] is a typical 3D model based method for pose invariant face recognition. The 3D morphable model is built using the principal component analysis of the 3D facial shapes and textures obtained from laser scanner devices. The 3D Morphable Model can then realize recognition either by transforming non-frontal face images to frontal view or by directly performing the recognition by using the coefficients of the morphable model. But usually, it is difficult to detect dense facial feature points that are accurate enough for the model fitting from low-resolution surveillance camera images.

Among the 2D appearance based methods, one of the successful approaches is to first generate a virtual frontal view by applying pose transformation on any given non-frontal face view. The View-Transition Model (VTM) [6] is a noteworthy work for pose transformation that can construct human appearance models for different poses which have proper texture information from a limited number of input images. The VTM method transforms views of an object between different poses by linear transformation of pixel values in images. For each pair of poses, a transformation matrix is calculated from image pairs of the poses of a large number of training data. The VTM was further extended to Local VTM (LVTM) in a patch-wise way [7] and it was shown that a more satisfactory face recognition result can be achieved using the virtual frontal face view generated by the local patch based LVTM than the original global patch based VTM. Both the VTM and the LVTM methods use a general training image dataset consisting of faces of a large number of individuals viewed from both frontal and various profile angles. The linear transformations learned from the training dataset are applied to the probe non-frontal face images, either in a global way as in the VTM or in a local patch based way as in the LVTM, to generate the counterpart virtual frontal face image that is then fed into a general traditional face recognition engine.

A more specific description of VTM and LVTM are as follows: Given a training multi-pose face image dataset $\Theta : \{\mathbf{Q}_\phi^1, \dots, \mathbf{Q}_\phi^N, \mathbf{Q}_{\theta_1}^1, \dots, \mathbf{Q}_{\theta_L}^1, \dots, \mathbf{Q}_{\theta_L}^N\}$, where N and L denote the number of training subjects and the number of profile poses, respectively. $\theta_l, (l = 1, \dots, L)$ are the degrees of pose rotation angle. When the degree equals zero, it becomes the frontal pose and is discriminatively denoted as ϕ . $\mathbf{Q}_\phi^n, (n = 1, \dots, N)$ represents the frontal face image for the n -th subject in the vector form which is a column vector that has pixel values of the image as its elements, and $\mathbf{Q}_{\theta_l}^n, (l = 1, \dots, L, n = 1, \dots, N)$ represents the non-frontal face image for the n -th subject with the pose rotation angle θ_l . For an input probe non-frontal face image $\mathbf{P}_{\theta_l}$, the purpose is to generate its virtual frontal image $\mathbf{P}_\phi$ using the linear transformation learned from the training dataset. The

VTM can be applied for virtual frontal face generation by one or any number of input images. However, in the interest of simplicity, we describe the frontal face generation algorithm for one input non-frontal face image only, and assume that the training dataset consists of frontal-nonfrontal face image pairs with one rotation degree θ only. The VTM calculates the linear transformation $\mathbf{T}$ beforehand using the training dataset by solving the following equation [6]:

$$\begin{bmatrix} \mathbf{Q}_\phi^1 & \dots & \mathbf{Q}_\phi^N \end{bmatrix} = \mathbf{T} \begin{bmatrix} \mathbf{Q}_\theta^1 & \dots & \mathbf{Q}_\theta^N \end{bmatrix} \quad (1)$$

Then the VTM generates $\mathbf{P}_\phi$, which is the virtual frontal face image for the probe image, from the input non-frontal probe face image $\mathbf{P}_\theta$ as follows:

$$\mathbf{P}_\phi = \mathbf{T} \mathbf{P}_\theta \quad (2)$$

Faces of two persons might have similar parts although these faces are not in total similar. Thus transforming the input face image using the information of the entire face image of other individuals might degrade the characteristics of the input individual's face. In order to solve this problem, the LVTM transforms face patches that are partial images of a face image for each location in the face image, instead of transforming directly the entire global face image. That is to say, LVTM achieves face pose transformation by synthesizing a face image from partial face image patches.

Let $\mathbf{q}_{\phi(x,y)}$ and $\mathbf{q}_{\theta(x,y)}$ represent face patches with patch center location at (x, y) of corresponding frontal and non-frontal global face image planes $\mathbf{Q}_\phi$ and $\mathbf{Q}_\theta$, respectively. The LVTM learns the location specific linear transforms $\mathbf{T}_{(x,y)}$ in a similar way with the VTM as follows,

$$\begin{bmatrix} \mathbf{q}_{\phi(x,y)}^1 & \dots & \mathbf{q}_{\phi(x,y)}^N \end{bmatrix} = \mathbf{T}_{(x,y)} \begin{bmatrix} \mathbf{q}_{\theta(x,y)}^1 & \dots & \mathbf{q}_{\theta(x,y)}^N \end{bmatrix} \quad (3)$$

It should be noted that the LVTM transforms each local area of an image while the VTM transforms the entire area of an image. Then the virtual frontal appearances for each local patches can be generated as follows:

$$\mathbf{p}_{\phi(x,y)} = \mathbf{T}_{(x,y)} \mathbf{p}_{\theta(x,y)} \quad (4)$$

After this, the LVTM synthesizes an output frontal face image $\mathbf{P}_\phi$ from all the transformed local patches $\mathbf{p}_{\phi(x,y)}$. The pixel values of regions where face patches are overlapped are calculated by averaging the pixel values of the overlapped patch, as illustrated in Fig. 1. Experimental results showed that the LVTM can achieve a higher recognition rate than that of using VTM for pose transformation [7].

This paper further extends the LVTM and presents a framework for face recognition across poses based on virtual frontal view generation using the LVTM with local patches clustering, which is denoted as c-LVTM hereafter. The proposed c-LVTM can describe the inherent transforming relationship between pixel values of patch pairs in a more precise way, thus more realistic virtual frontal face images can be generated and a higher recognition rate can be obtained. The experimental results on a real-world face dataset demonstrated the superiority of the proposed method. The rest of this paper is organized as follows: Section 2 describes

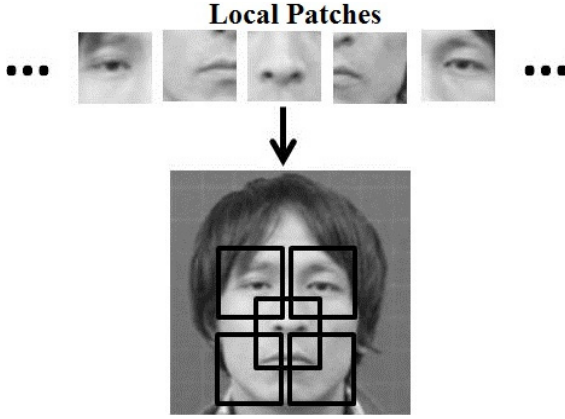


Fig. 1 Face image synthesis by local patches aggregation.

the proposed clustering based local VTM method (c-LVTM) in detail. Section 3 introduces the experimental result and Sect. 4 is the summary.

2. Virtual View Generation Using Clustering Based LVTM (c-LVTM)

The key point of VTM-like methods is the underlying linear relationship in the frontal and non-frontal face image pairs. Next we will show that the accuracy of the appearance transition model and the recognition rate can be further improved by better exploiting the inherent linear relationship between frontal-nonfrontal face image patch pairs. Based on the theoretical analysis and conclusion drawn in the reference literature [9] that – the more similar the 3D geometry of two objects are, the more similar the linear mapping is, we assume that for frontal-nonfrontal face image pairs, those patches with similar underlying 3D shapes and thus similar 2D appearances should have a more precise linear mapping relationship, since intuitively the variations in appearance caused by pose are closely related to the corresponding 3D structure. That is to say, frontal-nonfrontal patch pairs from more similar local 3D face structures should have a stronger linear relationship. Thus for the purpose of describing the relationship of frontal-nonfrontal pairs more precisely, it is better to learn the transformations specific to the local 3D structure. For each specific location, instead of learning a common transformation as in the LVTM, in the proposed c-LVTM, the corresponding local patches are first clustered based on the appearance similarity distance metric and then the transition models are learned separately for each cluster, rather than learning just a single common linear mapping using all the patch pairs for a specific location. As Fig. 2 (a) shows, in order to learn the linear transformations in a precise way, the training face image pairs are finely affine aligned using multiple landmarks. However in the testing stage, as Fig. 2 (b) shows, the input probe face image is only roughly affine aligned using three landmarks (left eye, right eye and nose tip), which can be easily detected by any off-the-shelf facial point detector.

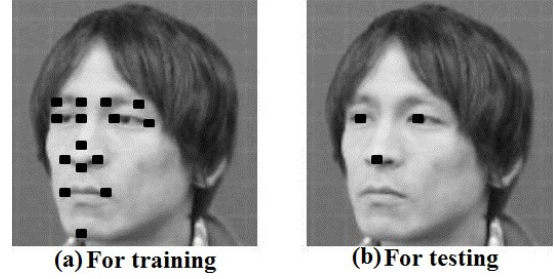


Fig. 2 Affine alignment using landmarks. Different strategies are used for training and testing stages. (a) In the training stage, in order to learn the linear transformations more accurately, the face images are finely affine aligned using multiple (15) landmarks labeled manually. (b) While in the testing stage, the input probe face image is only roughly affine aligned using three landmarks (left eye, right eye, and nose tip), which can be easily detected by any off-the-shelf facial point detectors.

More specifically, we first cluster the local patches $\mathbf{q}_{\theta(x,y)}$ for each location (x, y) into K clusters based on the appearance similarity where cluster k has c_k samples as $\{\mathbf{q}_{\theta(x,y)}^1, \dots, \mathbf{q}_{\theta(x,y)}^{c_k}\}$. Then for each cluster, the corresponding linear transformation $\mathbf{T}_{(x,y)}^k$, which is both location specific and local 3D structure specific, is learned as follows,

$$\begin{bmatrix} \mathbf{q}_{\phi(x,y)}^1 & \dots & \mathbf{q}_{\phi(x,y)}^{c_k} \end{bmatrix} = \mathbf{T}_{(x,y)}^k \begin{bmatrix} \mathbf{q}_{\theta(x,y)}^1 & \dots & \mathbf{q}_{\theta(x,y)}^{c_k} \end{bmatrix}, \quad (k = 1, \dots, K) \quad (5)$$

In the testing stage, the probe non-frontal face image is first roughly affine aligned, for example using only three landmark points at left eye, right eye and nose tip, which can be easily obtained using any standard facial feature point detector. Then for each local patch of the input non-frontal face image $\mathbf{p}_{\phi(x,y)}$, the most visually similar cluster in the training set is searched in the neighborhood regions $([x - \epsilon, x + \epsilon], [y - \epsilon, y + \epsilon])$ space of a specific location (x, y) . If we denote the most visually similar patch found resides in the k_{opt} -th cluster of location $(x_{\text{opt}}, y_{\text{opt}})$, then

$$\mathbf{p}_{\phi(x,y)} = \mathbf{T}_{(x_{\text{opt}}, y_{\text{opt}})}^{k_{\text{opt}}} \mathbf{p}_{\theta(x,y)} \quad (6)$$

The final transformed global frontal face image is aggregated from $\mathbf{p}_{\phi(x,y)}$ in a similar way as in the LVTM. The main idea of the appearance clustering based local transition models computation and the optimum transition model searching is illustrated in detail in Fig. 3. The differences between the VTM, the LVTM and the proposed c-LVTM are illustrated in Fig. 4. The VTM learns a global linear mapping on the holistic face image plane. The LVTM learns location specific linear mapping for each local patch. The proposed c-LVTM learns linear mappings that are both location specific and local 3D structure specific. As a whole, the flowchart of the proposed c-LVTM can be summarized as follows:

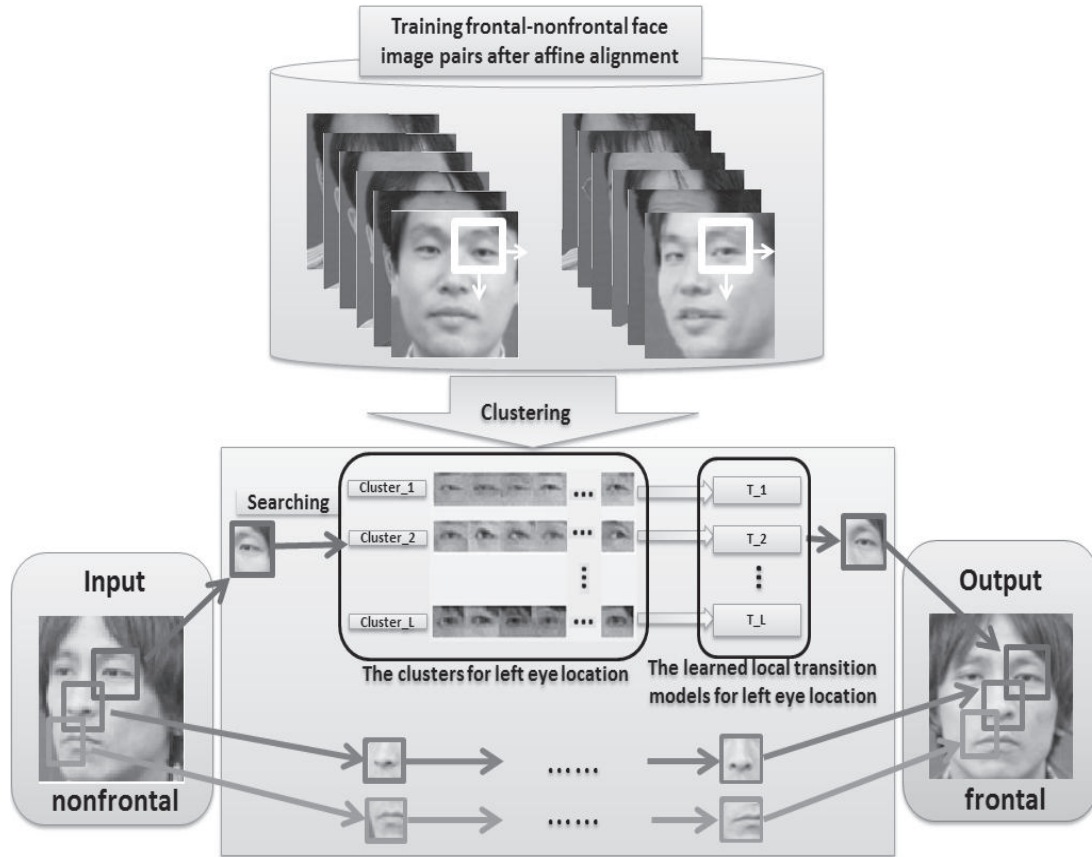


Fig. 3 An illustration of the main steps of the proposed c-LVTM method. The steps of the appearance clustering based local transition models computation and the optimum transition model searching are depicted by taking the local patches located on the left eye as an example. First, the local patches location on the left eye are clustered into clusters of **cluster_1**, **cluster_2**, ..., **cluster_L** based on appearance similarity. Then for each cluster, the local transition models **T_1**, **T_2**, ..., **T_L** are computed using the corresponding local patches. Then for the left eye local patch of an input non-frontal face image, the most visually similar clusters in the training set is searched in the neighborhood regions and local transition model corresponding to the most visually similar patch found is used to perform the transformation. The final transformed global frontal face image is the aggregation of all transformed local patches where the pixel values of the overlapped patches are averaged.

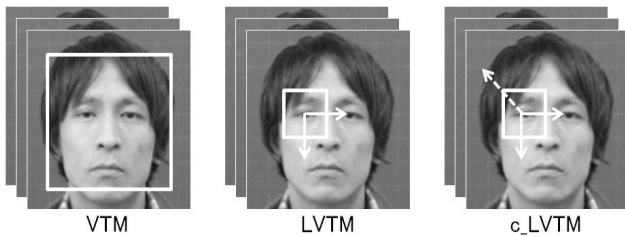


Fig. 4 The difference in how to select the image patterns for the transition model computation between VTM, LVTM and the proposed c-LVTM method.

The flowchart of the c-LVTM algorithm:

Training stage:

- 1) Finely affine align using multiple pre-labeled landmarks;
- 2) Split the face image plane into patches, and for each

patch location, perform clustering using appearance similarity;

- 3) For each cluster, compute the corresponding view transition model.

Testing stage:

- 1) Roughly affine align the given input probe non-frontal face image using auto-detected landmarks;
- 2) Split the face image plane into patches and for each patch search the most visually similar cluster in neighborhood range;
- 3) For each patch, using the searched optimum transition model, transform it into its frontal counterpart and aggregate all transformed patches;
- 4) Feed the transformed virtual frontal view face image into a general face recognition system.

It should be noted that the linear mapping is dependent on the underlying 3D structure and the way we try to find

the similar local structure is to cluster the corresponding 2D appearances. Because the 2D appearances are determined not only by the 3D structure but also by the corresponding textures, the proposed method is only an approximation of 3D structure clustering. Fortunately, usually the textures for the near facial locations are similar and do not change drastically. The proposed scheme based on the approximation assumption can still achieve a satisfactory result, which is also examined in the experiment.

3. Experiment

We used a subset of the face image dataset provided by SOFTPIA JAPAN to demonstrate the effectiveness of the

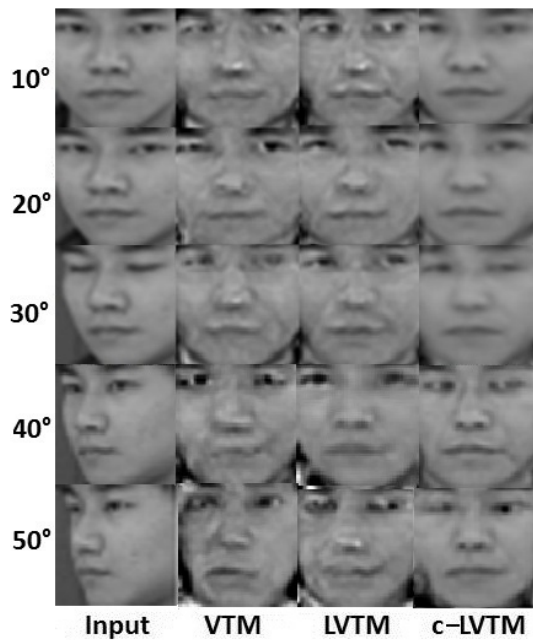


Fig. 5 The comparison of the visual effect of the transformed virtual frontal face image using different methods. It can be clearly seen that the virtual frontal face images generated using the proposed c-LVTM method have the highest visual fidelity.

proposed method. The subset consists of 250 individuals' images. They were taken with horizontal angles varying from -50 degrees to 50 degrees at 10 degrees interval. We compared the performance of using input images directly, the VTM, the LVTM and the proposed c-LVTM by 5-fold cross-validation. We transformed non-frontal face images to virtual frontal face images and then input the transformed images to a system that recognizes persons from the virtual frontal face images using a subspace based face recognition algorithm. The training images were precisely affine aligned using 15 landmark points and the testing images were roughly aligned using only 3 landmark points at left eye, right eye and nose-tip.

The image size was down-sampled to 32×32 in pixels, which is the usual size for detected face regions in low-resolution surveillance video. The face patch size was set to be 16×16 in pixels. The number of the cluster centers K was set to 4 . The region of neighborhood searching ϵ was set to 5 . The experiment was performed on a 2.8 GHz PC with 2 GB RAM memory under Matlab 7.12 platform. The comparison of computation time for training and testing stage on the dataset is illustrated in Table 1. The visual effects of the transformed virtual frontal face images using different methods are illustrated in Fig. 5. It should be noted that both the proposed method and its predecessors (L)VTM are first learning linear mappings over a relatively large training dataset and then applying the linear mapping to the input image, thus some subtleties in the output image might be influenced or biased by the training data. It can be seen that the generated virtual frontal face image using the proposed c-LVTM method has higher fidelity than that of other methods. This trend is further demonstrated by face recognition rate comparison which is illustrated in Fig. 6. The

Table 1 The comparison of computation time.

Computation time	VTM	LVTM	c-LVTM
Training stage	2 s	5 m	2.5 h
Testing stage	0.5 s	0.5 s	3 m

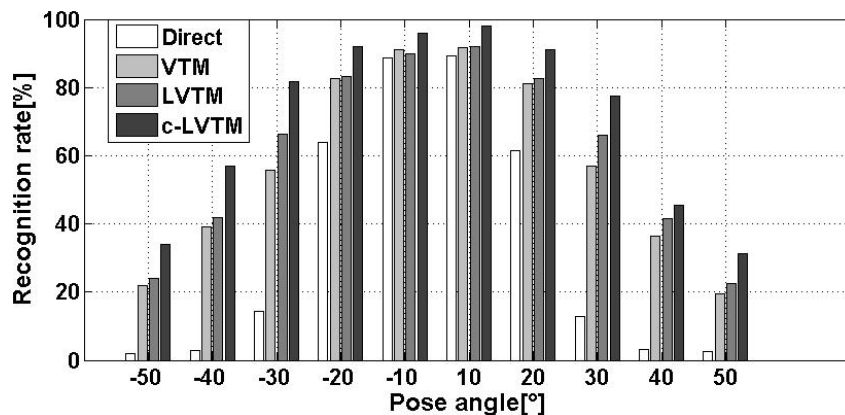


Fig. 6 Comparison of recognition rates across different angles. The input non-frontal face images are transformed using the VTM, LVTM and the proposed c-LVTM, respectively. The rate for the straightforward method of using the input non-frontal face images directly is also included for comparison.

Table 2 The comparison of c-LVTM and 3D morphable model.

Property	c-LVTM	3D morphable model
Computation Complexity	Intermediate	High
Pose Invariant Robustness	High	High
Facial Landmarks Needed	Auto	Manually
Lighting Invariant Robustness	Low	High
Special Device Needed	No	3D laser scanner

recognition rate of the straightforward baseline method that using the non-frontal face images directly as input is much lower than that of using the virtually generated frontal face images as input, either using VTM, LVTM or the proposed c-LVTM. Furthermore, the recognition performance of the proposed c-LVTM outperforms the VTM and LVTM in two ways: 1) c-LVTM has a higher recognition rate than VTM and LVTM; 2) Though all methods have a rate decreasing trend as the pose angle increases, the proposed c-LVTM has a more robust property against pose angle degree. That is to say, as the pose angle increases, the curve of rate-vs.-angle for c-LVTM drops less drastically than that of VTM and LVTM. The recognition rate comparison results validate our assumption that learning both location specific and local 3D structure specific linear transforms can better capture the relationship between frontal and non-frontal patch pairs than just learning a single common linear transformation. Table 2 is the property comparison of the proposed clustering based LVTM and the typical 3D morphable model method [4], where their corresponding pros and cons are illustrated in details.

4. Conclusion

In order to better exploit the underlying linear relationship between frontal and non-frontal pairs, this paper presented a framework for face recognition across pose based on virtual frontal view generation using the Local View Transition Model (LVTM) with local patches clustering. The proposed method further extended the LVTM by learning not only the local patch position specific transformations, but also the local 3D structure specific linear transforms. Experimental results showed the effectiveness of the proposed method.

We would like to further investigate in the following aspects: 1) The proposed method has the visual clustering and appearance searching procedures, that bring high computation burden. Next, we will study how to optimize the computation procedures to speed-up the computation more efficiently; 2) We plan to study the robustness of the proposed method to some specific uncertainties, for example, the subtle facial appearance changing and the error of the feature points detection; 3) We will study how to further improve the recognition rate of the proposed method, especially under the condition of larger profile degree.

Acknowledgements

This work was supported by “R&D Program for Implementation of Anti-Crime and Anti-Terrorism Technologies

for a Safe and Secure Society” Special Coordination Fund for Promoting Science and Technology of the Ministry of Education, Culture, Sports, Science and Technology, the Japanese Government. The face image dataset used in this work was provided by SOFTPIA JAPAN.

References

- [1] W. Zhao, R. Chellappa, P.J. Philips, and A. Rosenfeld, “Face recognition: A literature survey,” *ACM Computer Survey*, vol.35, no.4, pp.399–459, 2003.
- [2] M.A. Turk and A.P. Pentland, “Face recognition using eigenfaces,” *Proc. IEEE Computer Society Conference on Computer Vision and Pattern Recognition*, pp.586–591, 1991.
- [3] P.N. Belhumeur, J.P. Hespanha, and D.J. Kriegman, “Eigenfaces vs. Fisherfaces: Recognition using class specific linear projection,” *IEEE Trans. Pattern Anal. Mach. Intell.*, vol.19, no.7, pp.711–720, 1997.
- [4] V.G. Blanz, P.J. Phillips, and T. Vetter, “Face recognition based on frontal views generated from non-frontal images,” *Proc. IEEE Computer Society Conference on Computer Vision and Pattern Recognition*, pp.454–461, 2005.
- [5] D. Beymer, “Face recognition under varying pose,” *Proc. IEEE Computer Society Conference on Computer Vision and Pattern Recognition*, pp.756–761, 1994.
- [6] A. Utsumi and N. Tetsutani, “Adaptation of appearance model for human tracking using geometrical pixel value distribution,” *Proc. 6th Asian Conference on Computer Vision*, pp.794–799, 2004.
- [7] Y. Kono, T. Takahashi, D. Deguchi, I. Ide, and H. Murase, “Frontal face generation from multiple low-resolution non-frontal faces for face recognition,” *Proc. 10th Asian Conference on Computer Vision*, pp.175–183, 2010.
- [8] S. Baker and T. Kanade, “Hallucinating faces,” *Proc. IEEE Conference on Automatic Face and Gesture Recognition*, pp.83–88, 2000.
- [9] X. Chai, S. Shan, X. Chen, and W. Gao, “Locally linear regression for pose-invariant face recognition,” *IEEE Trans. Image Process.*, vol.16, no.7, pp.1716–1725, 2007.
- [10] D. Beymer and T. Poggio, “Face recognition from one example view,” *Proc. 5th International Conference on Computer Vision*, pp.500–507, 1995.



Xi Li received his B.S. degree in Electrical Engineering, and M.S. and Ph.D. degrees in Computer Science from Xi'an Jiaotong University, China. He is currently a postdoctoral researcher in the Graduate School of Information Science at Nagoya University, Japan. His research interests include pattern recognition and computer vision.



Tomokazu Takahashi received his B.S. degree from the Department of Information Engineering at Ibaraki University, and a M.S. and a Ph.D. from the Graduate School of Science and Engineering at Ibaraki University, Japan, in 1997, 2000, and 2003, respectively. He is currently an Associate Professor at Faculty of Economics and Information, Gifu Shotoku Gakuen University, Japan. His research interests include computer graphics and image recognition.



Daisuke Deguchi received B.S. and M.S. degrees in Engineering, and a Ph.D. degree in Information Science from Nagoya University, Japan, in 2001, 2003, and 2006, respectively. He is currently an Associate Professor at Information and Communications Headquarters, Nagoya University, Japan. He is working on the object detection, segmentation, recognition from videos, and their applications to ITS technologies, such as the detection and the recognition of traffic signs.



Ichiro Ide received his B.S. degree from the Department of Electronic Engineering, a M.Eng. degree from the Department of Information Engineering, and a Ph.D. from the Department of Electrical Engineering at the University of Tokyo, Japan, in 1994, 1996, and 2000, respectively. He is currently an Associate Professor in the Graduate School of Information Science at Nagoya University, Japan. His research interests include integrated media processing and video processing.



Hiroshi Murase received his B.S., M.S. and Ph.D. degrees from the Graduate School of Electrical Engineering at Nagoya University, Japan, in 1978, 1980, and 1987, respectively. He is currently a Professor in the Graduate School of Information Science at Nagoya University. He received the Ministry Award from the Ministry of Education, Culture, Sports, Science and Technology in Japan in 2003. He is a Fellow of the IEEE. His research interests include image recognition, intelligent vehicle, and com-

puter vision.