

Object Tracking using 2DLPP Manifold Learning

Huanlong Zhang^{1,2}, Shiqiang Hu¹, Lingkun Luo¹, Xiaolu Ke³

¹School of Aeronautics and Astronautics, Shanghai Jiao Tong University, Shanghai, 200240, China

²Department of Computer and Information Engineering, Luoyang Institute of Science and Technology, Henan, 471023, China

³Department of Automation, University of Science and Technology of China, Hefei, 230026, China
{Zh1_sjtu@126.com, sqhu@sjtu.edu.cn, lolinkun1988@sjtu.edu.cn, kxlu@mail.ustc.edu.cn}

Abstract—The task of visual tracking is to deal with dynamic image sequence. Traditional object representation in tracking algorithms using the image-as-vector subspace learning are easy to result in the problem of the curse of dimensionality and the loss of local structural information from the original image. In this paper, we present a novel online object tracking algorithm by using 2DLPP (Two-Dimensional Local Preserving Projections) manifold learning model. The proposed 2DLPP algorithm adopts a low dimensional eigenspace representation to reflect appearance changes of the target. It can preserve local structural information and directly extract features from image matrices, thereby the method facilitates the tracking task. Furthermore, the new method can update the feature basis recursively, and the computation becomes more efficient for online manifold learning of dynamic object. Finally, we apply the 2DLPP method to visual tracking in the particle filter framework. Experiment results demonstrate the effectiveness of the proposed method in different image sequences where the object undergoes large pose, scale, and lighting changes.

Keywords—2DLPP Tracking; Manifold Learning; Appearance Model; Nonlinear changes

I. INTRODUCTION

Object tracking is one of the most important components in a wide range of application in computer vision, such as motion perception, scene understanding and intelligent surveillance. The goal of visual tracking is to automatically locate the same object in an adjacent image frames from a video sequence once it is initialized. Despite numerous algorithms have been proposed in recent years, object tracking still remains many challenging problems due to appearance changes from pose, illumination, occlusion, and cluttered background.

To overcome the above problems, a variety of tracking algorithms have been proposed. These methods can be investigated in the following main aspects: (1) Generative tracking methods (e.g., [1, 2]), which drive the localization procedure using maximum-likelihood (ML) or maximum-a-posterior (MAP) formulation relying on target appearance model; (2) Discriminative tracking methods (e.g., [3-5]), which locate the target using a classifier that learns a decision boundary between the appearance of the target and that of the background; (3) Hybrid generative-discriminative tracking methods (e.g., [6, 7]), which combine the benefits of both the generative and the discriminative models. Among those methods, the trackers based on subspace representation in generative tracking model can maintain holistic appearance

information, which are more robust to in-plane rotation, scale change, illumination and pose changes.

Most of the earlier works of subspace models represent an image by a vector in high-dimensional space. The paper [8] proposed an eigen-tracking algorithm, which used the view-based eigenbasis representation for appearance changes. The paper [9] introduced Rao-Blackwellized particle filter for eigen-tracking. But, these methods are limited in real world due to the subspace constancy assumption. Recently there has been a trend of introducing machine learning techniques into tracking problem. The paper [10] proposed object tracking algorithm through incremental updating of the mean and eigenbasis for subspace learning, which has been one of the state-of-the-art methods in generative tracking framework. However, in real-world applications an image is intrinsically a matrix or a second order tensor. The above methods often confront the small-sample-size problem, where the number of samples is far smaller than the dimension of samples. Meanwhile, these methods based on the incremental image-as-vector subspace learning may lose local spatial information. To adapt these problems, the paper [11] proposed an tracking algorithm using incremental 2DPCA (two-dimensional Principal Component Analysis) learning and ML estimation, which develops the image-as-matrix learning algorithm for effective object tracking based on subspace model. And then, the paper [12] introduced L1-regularization into the 2DPCA reconstruction, which achieves more favorable performance than others. These methods do not need to turn image into high dimensional vector and reduce computational complexity of the covariance matrix. However, the efficiency of these methods heavily depended on the Euclidean structure. They fail to discover the underlying structure, if the images lie on a nonlinear submanifold hidden structure in the image space. Moreover, they can't preserve the local information of an image, which may be more important for object tracking.

In past few years, various locality preserving techniques including linear and nonlinear methods have been studied. Among them, the locality preserving projection (LPP) aims to preserve the intrinsic geometry of data and local structure. And experimental results demonstrate that LPP has better performance than PCA (Principal Component Analysis). Though LPP is linear projection, it is similar to nonlinear techniques such as Laplacian Eigenmaps^[13]. Recently, the LPP has been further researched. The paper [14] applies LPP to visual recognition and tracking, which shown it is effective for tracking under significant variations in pose, illumination and partial occlusion. However, the tracking algorithm based on the image-as-vector subspace method confronts the same

problem like 1D-PCA. Motivated by the idea of 2DPCA subspace method, the two-dimensional locality preserving projection (2DLPP) was proposed and applied to face recognition^[15] and Palmprint recognition^[16] successfully. The 2DLPP directly extracts feature and preserves local information, which has lower reconstruction errors.

Summarizing the above analysis, to the best of my knowledge, we firstly use 2DLPP method to the tracking problem and propose a novel subspace tracking algorithm (called 2DLPPT). The method could learn dynamic appearance changes of tracked object effectively in low dimensional subspace. When newly added image features arrive, the projection matrices can be obtained according to minimizing the objective function to keep the local structure of target appearance, and then the feature basis can also be updated in time. Meanwhile, in the particle filter framework visual tracking takes on the better performance. In contrast with the tracking algorithms using PCA or 2DPCA model, 2DLPPT can handle the nonlinear appearance changes and preserve local properties in high dimensional space. Experimental results on challenging image sequences demonstrate 2DLPPT is superior to other methods when the object undergoes the large pose, scale, and illumination changes.

II. PRELIMINARIES

Locality preserving projection (LPP) builds a graph incorporating neighborhood information of the data set. The objective of LPP is to preserve the local structure of the image space by explicitly considering the manifold structure. However, there are several disadvantages in application: (1) 2D image matrices must be previously transformed into 1D image vectors. The resulting image vector usually leads to a high dimensional image vector space, which can result in the curse of dimensionality problem. This problem is more apparent in small-sample-size cases; (2) Such a matrix-to-vector transform may cause the loss of structural information from the original image. To avoid these problems, we try to use 2DLPP to represent object appearance for robust visual tracking. In this section, we first briefly review the conventional 2DLPP method.

The 2DLPP method^[17] is a two-dimensional expansion of the vector-based LPP. Given a set of N sample images $X_1, X_2, \dots, X_N$, X_i represents an $m \times n$ image matrix ($m, n = 32$). 2DLPP seeks to constitute a local manifold subspace spanned by a set of projections W . The projected feature $y_i = X_i W$, where W is the $n \times d$ projection matrix. Then the X_i is reduced to lower dimension $m \times d$. To obtain a set of good projection function W , the objective function of 2DLPP is defined as:

$$\frac{1}{2} \min \sum_{i < j} \|y_i - y_j\|^2 S_{ij} \quad (1)$$

where y_i is the map of X_i . Minimizing the objective function is an attempt to ensure that if X_i and X_j are “close” the y_i and y_j are close as well. The S is similarity matrix, and $\|\cdot\|$ denotes the L_2 norm. There are several methods to measure “close”. Here, we adopt k -nearest neighbors to construct the nearest-neighbor graph. Meanwhile we use heat kernel to

determine the similarity S_{ij} . The detail elements in S is defined as follows:

$$S_{ij} = \begin{cases} \exp(-\|X_i - X_j\|^2 / \sigma), & \text{if } X_i \text{ is among } k \text{ nearest} \\ & \text{neighbors of } X_i \text{ or } X_j \text{ is among } k \text{ nearest} \\ 0 & \text{otherwise} \end{cases} \quad (2)$$

Here, k defines the local neighborhood, σ is a constant. The objective function with the selection of symmetric weights S_{ij} incurs a heavy penalty if $\|y_i - y_j\|^2$ is large. By simple algebra operation, we see that

$$\begin{aligned} \frac{1}{2} \sum_{i < j} \|y_i - y_j\|^2 S_{ij} &= \frac{1}{2} \sum_{i < j} \|(X_i - X_j)W\|^2 S_{ij} \\ &= \frac{1}{2} W^T \left(\sum_{i,j} (X_i^T X_i - X_i^T X_j) S_{ij} \right) W \\ &= \frac{1}{2} W^T \left(\sum_{i,j} X_i^T S_{ij} I_m X_i - \sum_{i,j} X_i^T S_{ij} I_m X_j \right) W \\ &= W^T X^T (L \otimes I_m) X W \end{aligned} \quad (3)$$

where $X = [X_1^T, X_2^T, \dots, X_N^T]^T$, D is a diagonal matrix whose entries are column (or row) sums of S . $L = D - S$ is the Laplacian matrix. $\otimes$ is the Kronecker product and I_m is an $m \times m$ identity matrix. Given the constraint is described as formula (4).

$$W^T X^T (D \otimes I_m) X W = 1 \quad (4)$$

The minimization problem of the objective function becomes

$$\begin{aligned} \min_W W^T X^T (L \otimes I_m) X W \\ \text{s.t. } W^T X^T (D \otimes I_m) X W = I_{d \times d} \end{aligned} \quad (5)$$

Using Lagrange multiplier and according to the KKT condition, the formula (6) is obtained.

$$X^T (L \otimes I_m) X W = \lambda X^T (D \otimes I_m) X W \quad (6)$$

Through solving the formula(5), the best projection matrix W can be obtained. 2DLPP has faster speed and higher recognition rate than LPP, and also has wider application rang than LPP when the number of training images is insufficient.

III. OBJECT TRACKING BASED ON MANIFOLD LEARNING

A. Particle Filter Framework

The basic idea of the particle filter (PF) is to use a set of weighted particles to approximate the true filtering distribution. The PF has achieved considerable success in visual tracking and became a widely used framework^[18,19]. In this paper, we design a novel tracking algorithm in particle filter framework, which gives an important tool for estimating the target of next frame without the concrete observation probability. According to given all available observations the posterior distribute of the state is predicted as follows:

$$p(x_t | y_{1:t-1}) = \int p(x_t | x_{t-1}) p(x_{t-1} | y_{1:t-1}) dx_{t-1} \quad (7)$$

where x_t describes the state of the target at the frames t . $y_{1:t-1} = \{y_1, y_2, \dots, y_{t-1}\}$ denotes the observation of the target from the first frame to the frames $t-1$. The $p(x_t | x_{t-1})$ represents the motion model between two consecutive states. The optimal state in frames t is obtained according to the MAP $x_t^* = \arg \max_x p(x | y_{1:t})$. Meanwhile, the state vectors are updated using Bayes rule, as shown in formula (8), where $p(x_t | y_t)$ denotes the likelihood function which is usually called observation model.

$$p(x_t | y_{1:t}) = \frac{p(y_t | x_t) p(x_{t-1} | y_{1:t-1})}{p(y_t | y_{1:t-1})} \quad (8)$$

In the particle filter framework, the posterior $p(x_t | y_t)$ is approximated by a finite set of n samples $\{x_t^i\}_{i=1}^n$ (called particles) with importance weights w_t^i . The particle samples x_t^i are drawn from an important distribution $q(x_t | x_{t-1}, y_{1:t})$ and the importance weights are updated as follows:

$$w_t^i = w_{t-1}^i \frac{p(y_t | x_t^i) p(x_{t-1}^i | x_{t-2}^i)}{q(x_t | x_{t-1}, y_{1:t})} \quad (9)$$

To avoid degeneracy, particles are resampled according to the importance weights so as to generate a set of equally weighted particles.

B. Motion Model

The six parameters of the affine transform are used to model $p(x_t | x_{t-1})$. Let the state $x_t = \{p_t, q_t, \theta_t, s_t, \alpha_t, \phi_t\}$, where $p_t, q_t, \theta_t, s_t, \alpha_t, \phi_t$ denote x, y translations, rotation angle, scale, aspect ratio, and skew respectively. Its matrix form H_A can be used to locate the patch from image like formula (10).

$$\hat{x} = H_A x = \begin{pmatrix} A & t \\ 0^T & 1 \end{pmatrix} x \quad (10)$$

where x and $\hat{x}$ are coordinates for pixels before and after affine transformation respectively. H_A is the affine transformation matrix which can be totally determined by the object state X_t . A is a 2×2 matrix as is shown formula (11). R is a rotation matrix which can be calculated by Rodrigues transform on an angle value and D is a diagonal matrix shown in (12).

$$A = R(\theta) R(-\phi) D R(\phi) \quad (11)$$

$$D = \begin{pmatrix} S & 0 \\ 0 & S \end{pmatrix} \begin{pmatrix} a & 0 \\ 0 & 1 \end{pmatrix} \quad (12)$$

The observation at frames t is the image patch located by its corresponding state X_t and it is denoted by Y_t . In tracking framework, the goal of object tracking is to determine the

posterior probability $p(X_t | Y_t)$. The tracking results are determined by the observation model.

In particle filter framework, to estimate 6 parameters which are corresponding to the affine transformation from the previous frame to the current one, motion model first predicts the states of particles and then spread particles in state space by adding Gaussian noise on each state item. The state transition is formulated by a Gaussian perturbation as follows:

$$p(x_t | x_{t-1}) = G(x_{t-1}, \sigma^2) \quad (13)$$

where G represents the Gaussian distribution with mean x_{t-1} and variance σ^2 . If the prior knowledge about the object can be obtained, states can be predicted more accurately.

C. Observation Model

For each observed image corresponding to a predicted state, how to determine the real state according to the candidate states is a key task. The observed image y_i is generated from a subspace spanned. After we obtain the optimal basis, the likelihood can be measured by the reconstruction error.

$$p(y_i | x_i) = \exp(-\|y_i - y_i^*\|_2^2) \quad (14)$$

where i denotes the i -th sample of the state. y_i^* denotes the reconstruction matrix according to the previous basis function. The reconstruction error determines which candidate state is the nearest to the real state.

A summary of the proposed tracking algorithm

1. Locate the target object in the first frame manually.
 2. Initialize the eigenbasis to be empty. The effective number of observations so far is $n = 1$.
 3. Advance to the next frame. Get resampling particles $\{x_t^i\}_{i=1}^N$ from the transition model (13).
 4. For each particle, extract the corresponding window F_t^i from the current window. Resize F_t^i to 32×32 pixel. Compute its weight from (9), which is its likelihood under the observation model.
 5. Store the image window corresponding to the most likely particle according to formula (14). The basis function is updated every ξ th frame. It can be calculated as follows
 - 1) Get the objects for frames t to frames $t + \xi$, and given the training data X , the size is $32 \times 32 \times \xi$.
 - 2) According to k -nearest neighbors get the relation matrix G , the size is $\xi \times \xi$.
 - 3) Get S_{ij} from heat kernel (2) based on G .
 - 4) Get Laplacian matrix $L = D - S$, D is the diagonal matrix. $D_{ii} = \sum_j S_{ji}$, $j = 1, \dots, Z$.
 - 5) Computer $X^T D X$ and $X^T L X$, get the eigenvalue λ and the projection matrix W from (5). According to the contribution of λ , the d is determined.
 6. Go to step 3.
-

IV. EXPEREMENTAL RESULTS

In this section, to demonstrate the improved performance of the proposed algorithm, we compare our 2DLPP tracker with 3 latest state-of-the-art trackers on 5 challenging image sequences. We use the source codes by the authors with the same initialization, and empirically tune its parameters for best performance. The three trackers are the weighted multiple instance learning tracking (WMILT)^[20], incremental subspace learning visual tracking (IVT)^[10] using 1DPCA model, and visual tracking using $\ell 1$ minimization (L1T)^[21]. The image sequence Mhyang, David2, Sylvester and Dudek are available on the website <http://visualtracking.net>. In addition, the image sequence Dollar is provided by the author^[20]. For the tracking methods which use particle filter (i.e., L1T, IVT), we use 600 particles in all tests.

A. Quantitative Analysis

All the tested image sequences are gray-scale. To evaluate the performance of these tracking methods in the 5 challenging image sequences we use two criteria^[22]. One is the success rate. The success frame is indicated as the following. Given the tracked bounding box ROI_T and the ground truth bounding box ROI_G , the score is defined as

$$score = \frac{area(ROI_T \cap ROI_G)}{area(ROI_T \cup ROI_G)} \quad (15)$$

The tracking results in one frame is considered as a success when this score is above 0.5. The success rate is defined as the number of success frames divided by the total number of frames in a video sequence. The other is location error which is defined as the Euclidean distance from the detected object center to the ground truth center at each frame. The average success rates and the mean of the center location errors are summarized in TABLE I and TABLE II respectively, where Bold fonts indicate the best result while the italic fonts indicate the second best ones.

TBALE I. SUCCESS RATES(%)

Vedio	Frames	IVT	L1T	WMILT	2DLPPPT
Dollar	327	100	73.9	76.1	100
Sylvester	1345	<i>85.4</i>	36.1	74.3	89.6
Mhyang	1490	99.0	90.0	36.9	<i>94.8</i>
Dudek	1145	<i>95.1</i>	50.0	68.7	96.8
David2	537	<i>80.7</i>	13.8	63.3	93.7

TBALE II. AVERAGE CENTER LOCATION ERRORS

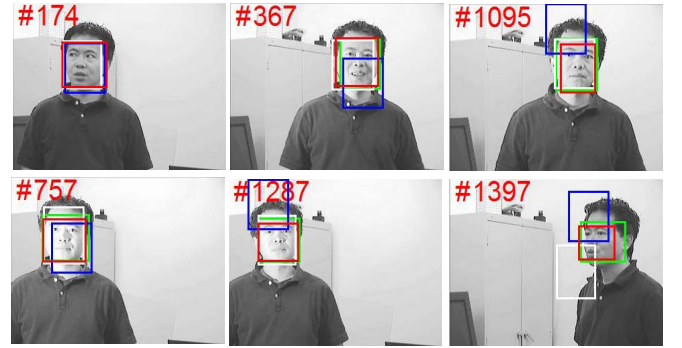
Vedio	Frames	IVT	L1T	WMILT	2DLPPPT
Dollar	327	8.2	22.4	14.4	4.5
Sylvester	1345	14.5	78.5	<i>11.4</i>	7.2
Mhyang	1490	2.6	9.8	28.1	<i>3.7</i>
Dudek	1145	10.5	90.8	35.6	<i>11.3</i>
David2	537	2.3	65.8	8.1	1.9

As shown in TABLE I and TABLE II, we can see our tracker achieves the best or the second best performance compared with IVT, L1T and WMILT. Though the average location error of our tracker is smaller than IVT on the image sequence Dudek, the success rates are just reverse. In general, our tracker and IVT are the better performance than others for

5 challenging image sequences with pose, scale and illumination changes.

B. Quantitative Analysis

For Mhyang and Dudek sequence shown in Figure 1(a) and (b), the illumination, scale and pose of the object change gradually. The IVT algorithm performs well on these sequences, as it is designed to account for affine motion and appearance changes. IVT adapts a low dimensional eigenspace representation to reflect appearance changes of the target based on PCA method, which focus on the global information. Under the condition of smooth appearance changes, our tracker has a little better than IVT on the dudek sequence, but IVT shows the better performance generally. The WMIL algorithm adopts refined Haar-like feature in boosting framework. When drastic illumination changes occur on Mhyang sequence (#757, #1287, #1095), WMIL fails to track object as the capability of the classifier is degraded. Moreover, L1T makes the candidate with the smallest projected error be the tracking object, which takes advantage of the original pixel. L1T can perform well on the Mhyang image sequence with a relative high resolution and smooth changes. But when object undergoes out-of-plane rotation on Dudek sequence (#571, #770, #865), L1T fails to track object as the template is hard to be updated accurately.



(a) Tracking results on Mhyang sequence



(b) Tracking results on Dudek sequence

— IVT — L1T — WMILT — 2DLPPPT

Fig. 1. The results of our tracker and others

For the David2 sequence shown in Figure 2, it takes on drastic pose changes (#172, #346, #421). Under the condition of this appearance changes, our tracker successfully tracks the right object on this sequence as it adopts 2DLPP which can keep local structure in high dimensional space. However, the IVT and WMILT algorithms show instability. Among them,

the IVT updates the eigenbasis in real time in order to represent the appearance changes. But too drastic pose changes make the basis function adapt it hardly resulting in tracking drift. The WMILT considers the background information so that it tracks object completely on this sequence with the simple environment. If the selected samples are not accurate the negative samples are easy to become the tracked object so that the WMILT algorithm tracks object inaccurately. The LIT fails to track object since the frames 74 as the poor image feature and large deformation.

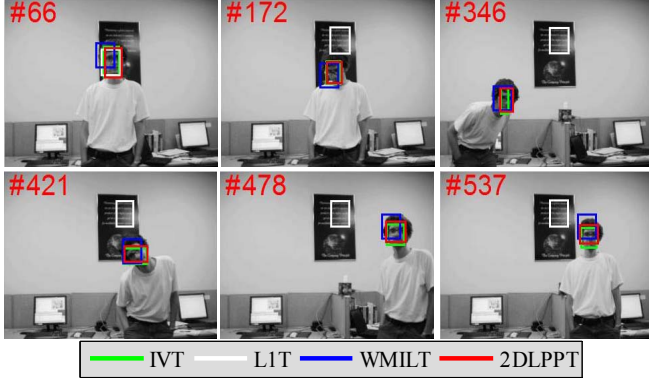


Fig.2. The results of our tracker and others

For the Sylvester sequence shown in Figure 3, it contains an animal doll moving in different pose, scale, and lighting conditions (#270,#694,#787,#959). Our tracker achieves a better performance than other methods. The LIT does not adapt the problem of out-of-plane rotation so as to tracking failure since frames 425. IVT also fails to track object when object encounters the larger pose change at 1180 frame. However, WMILT can track object completely as it extracts discriminative features online with classifier updating for foreground and background separation. But, in contrast of our tracker the tracked results of the WMILT method are not accurate enough.

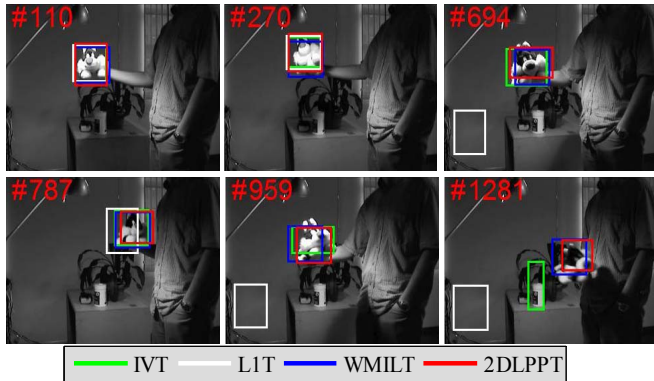


Fig.3. The results of our tracker and others

For the Dollar sequence shown in Figure 4, object undergoes significant appearance changes at the 60th frame and similar object distractor giving rise to background clutter. All the algorithms can track object completely, since the object appearance is nearly constant and motion is smooth. However, they show drift from the real trajectory. In contrast of them, our tracker gets the smallest location errors.

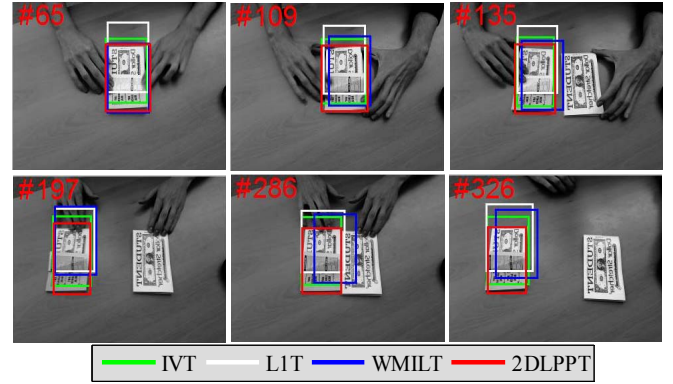


Fig.4. The results of our tracker and others

C. Quantitative Analysis

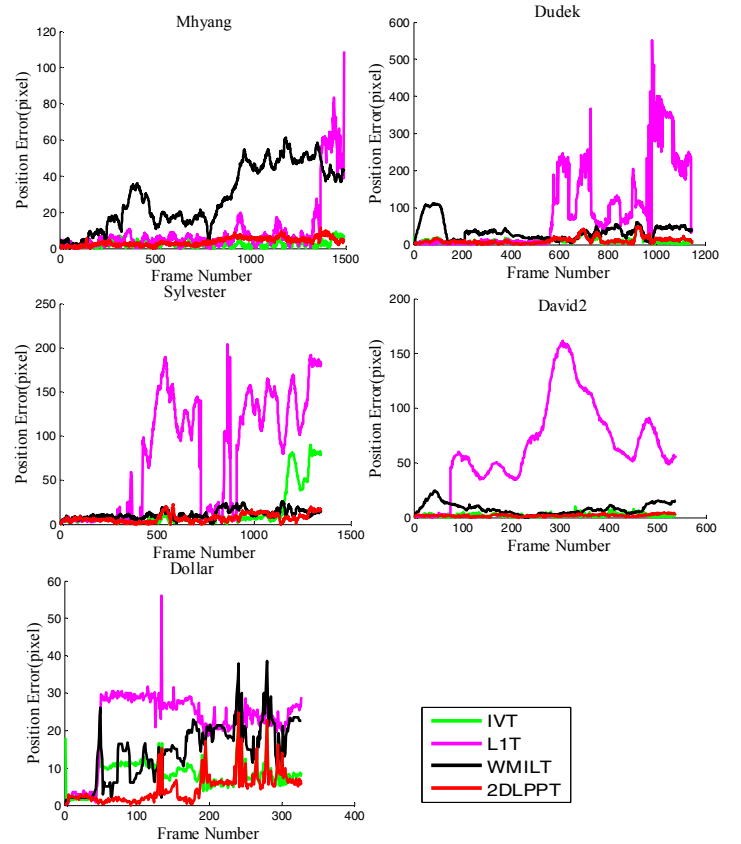


Fig.5. Error plots of all the tested sequence

For describing the tracking results in more detail, the Fig.5 shows the location errors changes using different algorithms in all the test sequences. Our tracker achieves the best performance for the most of sequences in terms of accuracy and robustness. Combined with TABLE I, TABLE II, Fig.1, Fig.2, Fig.3 and Fig.4, we can conclude that 2DLPPT algorithm can track the object favorably against other methods when object undergoes the pose, scale, and illumination changes, especially when the appearance changes of the tracked object are nonlinear.

V. CONCLUSION

We propose an appearance-based tracker that incrementally learn a low dimensional eigenspace representation for visual object tracking when the object undergoes large pose, illumination and scale changes. Whereas most existing tracking algorithms based on 1DPCA method and 2DPCA method fail to discover the local structural information and are hard to represent the nonlinear hidden feature. The trackers based on LPP model overcome the above problems to some extent. But they may face the problems of the curse of dimensionality and the high computing cost so as to process image sequences difficultly. We adopt 2DLPP manifold learning method to track visual object. The method can show the nonlinear appearance changes and keep the local structural information in low dimensional feature space. Moreover, by updating basis function a good eigenspace representation is obtained to model the appearance of the tracked object. In the particle framework, our method shows the better performance than others on some challenging image sequences.

Acknowledgment

This paper is supported by the National Natural Science Foundation of China (61074106, 61374161).

REFERENCES

- [1] J. Ning, L. Zhang, D. Zhang, and C. Wu. Scale and orientation adaptive mean shift tracking. *Computer Vision, IET*, 6, pp.52-61, 2012.
- [2] LongyinWen, ZhaoweiCai, Zhen Lei, Dong Yi, and Stan Z. Li. Online Spatio-temporal Structural Context Learning for Visual Tracking. *ECCV*. 75, pp.716-729, 2012.
- [3] Jialue Fan, Ying Wu, and Shengyang Dai. Discriminative Spatial Attention for Robust Tracking. *ECCV*, 6311, pp.480-493, 2010.
- [4] Nan Jiang, Wenyu Liu, Ying Wu. Learning Adaptive Metric for Robust Visual Tracking. *IEEE Trans. IMAGE PROCESSING*. 20 (8), pp.2288-2291, 2011.
- [5] Sam Hare, Amir Saffari, Philip H. S. Torr. Structured Output Tracking with Kernels. *ICCV*, pp.263-271, 2011.
- [6] Thang Ba Dinh, Gerard Medioni. Co-training Framework of Generative and Discriminative Trackers with Partial Occlusion Handling. pp. 642 – 649, 2011.
- [7] T. Woodley, B. Stenger, and R. Cipolla. Tracking using online feature selection and a local generative model. pp.1-8, 2007.
- [8] M. Black and A. Jepson. Eigenttracking: Robust matching and tracking of articulated objects using a view-based representation. *IJCV*, vol. 26(1), pp. 63-84, 1998.
- [9] Z. Khan, T. Balch, F. Dellaert. A rao-blackwellized particle filter for eigentracking. *CVPR*, vol.2, 2004.
- [10] D. A. Ross, J. Lim, R.-S. Lin, and M.-H. Yang. Incremental learning for robust visual tracking. *International Journal of Computer Vision*, vol. 77, pp. 125-141, 2008.
- [11] T. Wang, I. Y. H. Gu, and P. Shi. Object tracking using incremental 2d-pca learning and ml estimation. in *ICASSP*, pp.933–936, 2007.
- [12] Dong Wang and Huchuan Lu. Object Tracking via 2DPCA and L1-Regularization. *IEEE PROCESSING LETTERS*, vol.19(1), pp.711-715, 2012.
- [13] X.F. He, P. Niyogi. Locality preserving projections. in: *Proceedings of the Conference on Advances in Neural Information Processing Systems*, 2003.
- [14] Yanxia Jiang, Hongren Zhou, Zhongliang Jing. Visual tracing and recognition based on robust locality preserving projection. *Optical Engineering*, vol.46(4), 046401, 2007.
- [15] Ben Niu, Qiang Yang, Semon Chi Keung Shiu, et. Two-dimensional Laplacian faces method for face recognition. *Pattern Recognition*, 41, pp. 3237-3243, 2007.
- [16] Deweb Hu, Guiyu Feng, Zongtan Zhou. Two-dimensional locality preserving projections (2DLPP) with its application to palmprint recognition. *Pattern Recognition*. 40, pp.339-342, 2007.
- [17] Hao-Xin Zhao, Hong-Jie Xing, Xi-Zhao Wang, and Jun-Fen Chen. L1-Norm-Based 2DLPP. *Chinese Control and Decision Conference*, pp.1259-1264, 2011.
- [18] Yuru Wang, Xianglong Tang, Qing Cui. Dynamic appearance model for particle filter based visual tracking. *Pattern Recognition*. 45, pp.4510-4523, 2012.
- [19] Samarjit Das, Amit Kale, and Namrata Vaswani. Particles Filter With a Mode Tracker for Visual Tracking Across illumination Changes. *IEEE TRANSACTIONS ON IMAGE PROCESSING*, vol.21(2), pp.2340-2347, 2012.
- [20] Kaihua Zhang, Huihui Song. Real-time visual tracking via online weighted multiple instance learning. *Pattern Recognition*. 46, pp.397-411, 2013.
- [21] X. Mei and H. Ling. Robust visual tracking using ℓ_1 minimization. in *Computer Vision, 2009 IEEE 12th International Conference on*, pp. 1436-1443, 2009.
- [22] M. Everingham, L. Van Gool, C. K. Williams, J. Winn, and A. Zisserman. The pascal visual object classes (voc) challenge. *International journal of computer vision*, vol. 88, pp.303-338, 2010.