

A novel illumination-robust local descriptor based on sparse linear regression [☆]



Zuodong Yang ^{a,b}, Yong Wu ^{a,b}, Wenteng Zhao ^{a,b}, Yicong Zhou ^c, Zongqing Lu ^{a,b},
Weifeng Li ^{a,b,*}, Qingmin Liao ^{a,b}

^a Department of Electronic Engineering, Graduate School at Shenzhen, Tsinghua University, China

^b Shenzhen Key Laboratory of Information Science and Technology, Guangdong, China

^c Department of Computer and Information Science, University of Macau, Macau, China

ARTICLE INFO

Article history:

Available online 13 October 2015

Keywords:

Face recognition
Illumination-insensitive representation
Local descriptor
Sparse linear regression

ABSTRACT

Robust face recognition under uncontrolled illumination conditions is an important problem for real face recognition systems. In this paper, we introduce a novel illumination-robust local descriptor named Sparse Linear Regression Binary (SLRB) descriptor. The SLRB descriptor is a bit string by binarizing the sparse linear regression coefficients in a local block. It is an illumination-insensitive descriptor based on the locally linear consistency assumption under the Lambertian reflectance model. We use the cosine similarity and Hamming similarity as the similarity measure for the SLRB descriptor of two different images respectively. Experimental results on the Extended Yale-B and CMU-PIE face database show a promising performance compared to the existing representative approaches.

© 2015 Elsevier Inc. All rights reserved.

1. Introduction

Face recognition has recently received a lot of attention and has been applied to the fields of entertainment, smart cards, information security and law enforcement and surveillance [1]. Despite tremendous advance in face recognition has made, robust face recognition under uncontrolled illumination conditions is still challenging [2,3]. In recent years, there are a number of approaches for dealing with face image variations due to illumination changes. They can be roughly classified into four categories.

The first category attempts to handle the illumination normalization problem with traditional image processing methods. Histogram Equalization (HE) obtains an image with a high dynamic range and a great deal of details based on information in the histogram [4]. Gamma Intensity Correction (GIC) was introduced to normalize the overall image intensity at the given grey-level by Gamma transformation [5]. Logarithm Transform (LT) was proposed to perform illumination normalization in face im-

ages under uncontrolled illumination conditions [6]. These methods are mainly based on intensity transformation and used as pre-processing methods.

Using the illumination samples, the second category learns the model of face images under varying illumination. In [7] the author made the explanation that arbitrary illumination conditions could be modeled by an image basis and showed that five eigenfaces suffice to represent face images under a wide range of lighting conditions. In [8], it showed the fact that a set of images of an object, which has a fixed pose under varying illumination conditions, form a convex illumination cone in the space of images. In [9], it was proved that the intensity of the object surface obtained with arbitrary distance light sources spans a 9-dimension linear subspace based on a spherical harmonic representation. In [10], the authors propose a novel framework named Face Analysis for Commercial Entities (FACE) and adopt normalization (“correction”) strategies to address illumination variations. These methods depend on a statistical model or a physical model, and can settle the illumination variations well. However, they require a large amount of training samples under varying illumination conditions in most cases, which makes them not practical for real face recognition systems.

The third category deals with illumination variations by removing the illumination component. Jobson et al. introduced the Retinex approach to obtain the reflectance component by estimating the illumination component [11,12]. In [13] the author enhanced the illumination removal phase and used double-density dual-tree complex wavelet transform (DD-DTCWT) filtering

[☆] This work was supported in part by the National Natural Science Foundation of China (Grants No. 61271393 and No. 11002082).

* Corresponding author at: Department of Electronic Engineering, Graduate School at Shenzhen, Tsinghua University, China.

E-mail addresses: yangzd13@mails.tsinghua.edu.cn (Z. Yang), wuyong11@mails.tsinghua.edu.cn (Y. Wu), zhaowt13@mails.tsinghua.edu.cn (W. Zhao), yicongzhou@umac.mo (Y. Zhou), luzq@sz.tsinghua.edu.cn (Z. Lu), Li.Weifeng@sz.tsinghua.edu.cn (W. Li), liaoqm@tsinghua.edu.cn (Q. Liao).

to extract the reflectance portion. Some methods were proposed to remove the illumination component in transformation domain as well, such as the Homomorphic filtering approach [14], discrete cosine transform in the logarithm domain [15], wavelet transform in the frequency domain [16], etc. These methods employ the fact that illumination is a low frequency component and can be applied to a single image without many training samples.

The forth category attempts to find a representation which is insensitive to illumination variations. Wang et al. defined the ratio of the albedo of one image as the Self-Quotient Image (SQI) which is independent of illumination [17]. In [18] the authors proposed the Relative Image Gradient (RIG) feature which is robust against to illumination variations. Gradientfaces takes the similar idea but utilizes the ratio between x -gradient and y -gradient [19]. Both RIG and Gradientfaces are extracted in the gradient domain and can be applied to a single sample. In addition, in [20] and [21], the authors introduced the Weberface and Generalized Weberface (GWF) which employ the relative intensity difference between the center pixel and its neighborhoods derived from the Weber's law. Most of the above methods make use of the Lambertian reflectance model and are based on the assumption that the illumination component is characterized by slow variations while the reflectance component varies drastically.

Apart from the above methods that dedicate to illumination normalization, some approaches for texture classification have been employed for illumination-robust face recognition as well. Local Binary Pattern (LBP) [22] is one of the most commonly used methods. LBP is a local descriptor of texture which processes the difference between the intensity of the center pixel and its neighborhoods with binary encoding. It has been widely applied to illumination-robust face recognition due to its tolerance of monotonic illumination variations and computational simplicity. Nevertheless, LBP is sensitive to noise when the image region is near-uniform. To improve the noise robustness, Local Ternary Pattern (LTP) was proposed in [23], which utilized ternary encoding instead of binary encoding in LBP.

In this paper, we propose a novel local binary descriptor named the Sparse Linear Regression Binary (SLRB) descriptor based on the sparse linear regression. The SLRB descriptor is a bit string obtained by binarizing the sparse linear regression coefficients in a local patch. We prove the SLRB descriptor to be robust to illumination based on the following assumptions.

1) The intensity of the center pixel, $f(x_0, y_0)$, can be linearly expressed by these of its neighborhoods and the linear combination coefficients α_i are consistent in a local block, which is similar to the assumption used in Locally Linear Embedding (LLE) [24]:

$$f(x_0, y_0) = \sum_{j=1}^N \alpha_j f(x_j, y_j) + \epsilon, \quad (1)$$

where (x_0, y_0) is the coordinate of the center pixel, (x_j, y_j) is the coordinate of the surrounding pixel and ϵ is the residual term.

2) The intensity of an image, $f(x, y)$, can be expressed as the product of its illumination component, $i(x, y)$, and reflectance component, $r(x, y)$, which is indicated in [14]:

$$f(x, y) = i(x, y)r(x, y). \quad (2)$$

In addition, the illumination component varies slowly while the reflectance component changes abruptly.

We summarize the characteristics of the proposed SLRB descriptor as follows.

1) Based on the Lambertian reflection model, we can prove that the SLRB descriptor is an illumination-insensitive feature and can be effectively applied to illumination-robust face recognition compared with the LBP and LTP features.

2) The SLRB descriptor is a bit string with a low dimension and we can simply employ the cosine distance and Hamming distance as similarity measures. Therefore, the SLRB descriptor is quite efficient and requires less computation complexity than other methods.

3) The lasso regression [25] exhibits the stability of the ridge regression method. Binary encoding scheme is a common method in face recognition and can reduce local noise. As a result, the SLRB descriptor is robust to noise on account of lasso regression and binary encoding.

The rest of the paper is organized as follows. We present our SLRB descriptor, the illumination-robust face recognition algorithm based on the SLRB descriptor and prove its illumination robustness in the next section. In Section 3, we illustrate some experiments by applying our face recognition algorithm on the Extended Yale Face Database B and CMU-PIE face database. Finally, we conclude our paper in Section 4.

2. Sparse linear regression binary descriptor

2.1. Local descriptor based on linear regression

As mentioned above, suppose that the intensity of the center pixel is a linear combination of the intensity of its neighborhoods as Eq. (1). To obtain the linear combination coefficients, we make a further assumption that the linear combination coefficients are the same in a local block. In an $N \times N$ block, we suppose that the intensity of the center pixel is linearly expressed by that of the surrounding eight pixels in a patch. Then there are $(N-2)^2$ "center pixels" $f^{(k)}$, $k = 1, \dots, (N-2)^2$, that form a column vector, dependant variable $\mathbf{f} = [f^{(1)}, f^{(2)}, \dots, f^{((N-2)^2)}]^T$.

For each patch, we have

$$\begin{aligned} f^{(k)} &= \sum_{j=1}^8 \alpha_j f_j^{(k)} + \epsilon_k \\ &= \mathbf{x}_k^T \boldsymbol{\alpha} + \epsilon_k, \quad k = 1, 2, \dots, (N-2)^2, \end{aligned} \quad (3)$$

where $f_j^{(k)}$ is the intensity of the j -th surrounding pixel of the k -th center pixel, the surrounding pixel vector $\mathbf{x}_k = [f_1^{(k)}, f_2^{(k)}, \dots, f_8^{(k)}]^T$, and the linear regression coefficient vector $\boldsymbol{\alpha} = [\alpha_1, \alpha_2, \dots, \alpha_8]^T$.

These $(N-2)^2$ equations can be stacked together and written as

$$\begin{bmatrix} f^{(1)} \\ f^{(2)} \\ \vdots \\ f^{((N-2)^2)} \end{bmatrix} = \begin{bmatrix} \mathbf{x}_1^T \\ \mathbf{x}_2^T \\ \vdots \\ \mathbf{x}_{(N-2)^2}^T \end{bmatrix} \boldsymbol{\alpha} + \begin{bmatrix} \epsilon_1 \\ \epsilon_2 \\ \vdots \\ \epsilon_{(N-2)^2} \end{bmatrix}. \quad (4)$$

Eq. (4) can be reformulated as

$$\mathbf{f} = \mathbf{X}\boldsymbol{\alpha} + \boldsymbol{\epsilon}, \quad (5)$$

in which the surrounding pixel matrix $\mathbf{X} = [\mathbf{x}_1^T, \mathbf{x}_2^T, \dots, \mathbf{x}_{(N-2)^2}^T]^T$ and the residual vector $\boldsymbol{\epsilon} = [\epsilon_1, \epsilon_2, \dots, \epsilon_{(N-2)^2}]^T$.

By solving Eq. (5), the linear regression coefficient vector $\boldsymbol{\alpha}$ would be determined as a descriptor of the current block. Fig. 1 illustrates the procedures of the linear regression coefficients in a block.

2.2. Sparse coefficients and binary encoding

Generally, we can determine the linear regression coefficient vector $\boldsymbol{\alpha}$ by the ordinary least square (OLS) estimation:

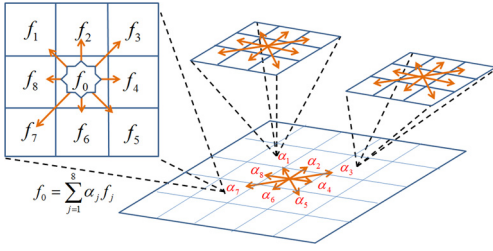


Fig. 1. Illustration of the SLRB descriptor in a block.

$$\underset{\alpha}{\operatorname{argmin}} \|\mathbf{f} - \mathbf{X}\alpha\|^2, \quad (6)$$

which has the solution

$$\hat{\alpha} = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{f}. \quad (7)$$

Or address this problem by the ridge regression method for the underdetermined system of equations as

$$\underset{\alpha}{\operatorname{argmin}} \|\mathbf{f} - \mathbf{X}\alpha\|^2 + \lambda_2 \|\alpha\|^2 \quad (8)$$

with the solution

$$\hat{\alpha} = (\mathbf{X}^T \mathbf{X} + \lambda_2 \mathbf{I})^{-1} \mathbf{X}^T \mathbf{f}, \quad (9)$$

where λ_2 is the ℓ_2 -norm regularization parameter, and $\mathbf{I}$ is the identity matrix.

Furthermore, as [25] indicated, the lasso regression has a further advantage over the ridge regression method of producing interpretable submodels and exhibits the stability of the ridge regression method. Therefore, we employ the lasso regression for the solution of α instead:

$$\underset{\alpha}{\operatorname{argmin}} \|\mathbf{f} - \mathbf{X}\alpha\|^2 + \lambda_1 \|\alpha\|_1, \quad (10)$$

where λ_1 is the ℓ_1 -norm regularization parameter. The lasso regression tends to produce some sparse coefficients that exactly are 0 and gives a more interpretable descriptor of the local texture.

Subsequently, we adopt binary encoding of the sparse regression coefficients for further dimension reduction as the Sparse Linear Regression Binary (SLRB) descriptor $\beta = [\beta_1, \beta_2, \dots, \beta_8]^T$ of the current image block:

$$\beta_j = \begin{cases} 1 & \text{if } \alpha_j > \mu_\alpha \\ 0 & \text{otherwise,} \end{cases} \quad (11)$$

where μ_α is the mean of α .

Fig. 2 shows a face block of 6×6 with non-uniform illumination. We can obtain the SLRB descriptor $\beta = [0, 0, 1, 0, 1, 0, 0, 1]^T$ and the L^2 -norm of the residual vector $\|\epsilon\| = 4.11$, which is much smaller than the intensity of the pixels and can be ignored. For every small block of 3×3 , we can get an SLRB descriptor and the mean value of these 16 SLRB descriptors is $\mu_\beta = [0, 0, 1, 0, 1, 0, 0.69, 1]^T$, which is almost the same as β and proves that our assumption is reasonable.

2.3. Illumination robustness

As mentioned above, it is assumed that the intensity is the product of its illumination component, which is a constant in a local area approximatively, and reflectance component that represents the face feature as Eq. (2). It is easy to prove that the proposed SLRB descriptor is an illumination insensitive representation as follows.

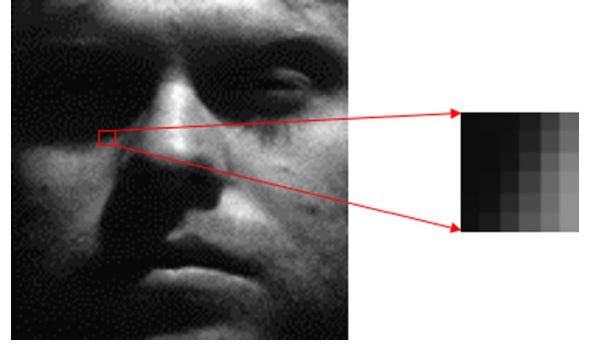


Fig. 2. A face block with non-uniform illumination.

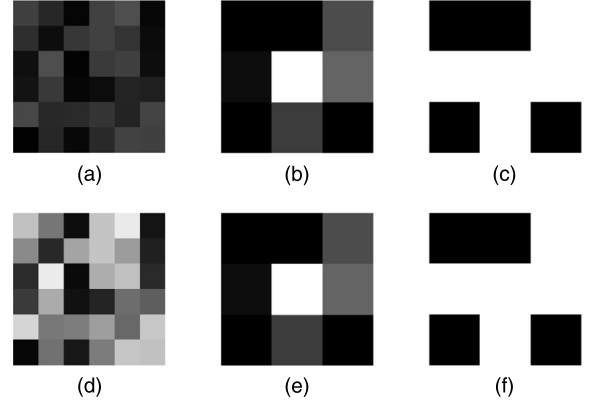


Fig. 3. The robustness to illumination variations of the SLRB descriptor under the Lambertian reflectance model. (a) and (d) are the same sample blocks under different illumination conditions; (b) and (e) are the linear regression coefficients of (a) and (d) respectively; (c) and (f) are the SLRB descriptors, i.e., binary encoding of (b) and (e).

Combining Eq. (1) and Eq. (2), we have

$$i(x_0, y_0)r(x_0, y_0) = \sum_{j=1}^N \alpha_j i(x_j, y_j)r(x_j, y_j) + \epsilon. \quad (12)$$

Based on the assumption that the illumination component varies slowly in a local area, i.e.

$$i(x_0, y_0) \approx i(x_j, y_j), \quad (13)$$

Eq. (12) can be reformulated as follows by ignoring the residual term.

$$r(x_0, y_0) \approx \sum_{j=1}^N \alpha_j r(x_j, y_j) \quad (14)$$

From Eq. (14), we prove that the linear regression coefficients are illumination insensitive representation. They depend only on the reflectance component and have nothing to do with the illumination component. Thus the SLRB descriptor which is the binary encoding of the linear regression coefficients is robust to illumination variations as well. Fig. 3 shows an example of the SLRB descriptor against illumination variations. The original image block (a) is produced randomly with the size of 6×6 and (d) is the same image block with different illumination components. It is obvious that their linear regression coefficients ((b), (e)) and SLRB descriptors ((c), (f)) are the same.

2.4. Computational efficiency

The SLRB descriptor is employed as an illumination-insensitive feature for face recognition. The face image is divided into some

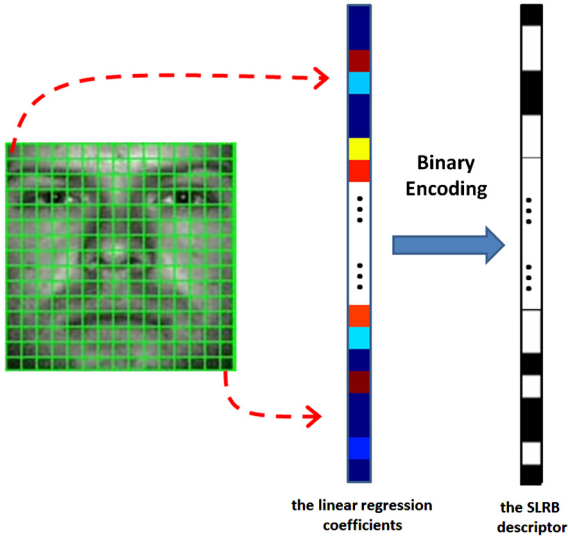


Fig. 4. Illustration of feature extraction. The linear regression coefficients are calculated in each block and then concatenated. The feature is a bit string by binarizing the regression coefficients.

blocks in which the SLRB descriptor is calculated. Denote $\beta^{(i)}$ is the SLRB descriptor of the i th block, then the feature of a face image can be written as $[\beta^{(1)}, \beta^{(2)}, \dots, \beta^{(M)}]^T$, where M is the number of blocks. Then the feature is extracted which is a bit string with a low dimension. Fig. 4 illustrates the extraction of the feature based on the SLRB descriptor in a face image. Take an 8-bit gray-scale image with the size of 120×120 as an example. The features of LBP and LTP are 13 275 bits and 115 200 bits respectively while the feature based on the SLRB descriptor is 7200 bits when the size of blocks is 4×4 and 1800 bits when the block size is 8×8 . They are 1/16 and 1/64 of the number of storage bits of the original image and much smaller than the features of LBP and LTP. Therefore, the SLRB descriptor is very efficient to compute and requires less computation complexity than other methods.

In addition, the SLRB image can be defined for a further illustration. The SLRB image is calculated by assigning a binomial factor 2^l for the SLRB descriptor of the block centered on the current pixel, which is similar to the LBP image. Fig. 5 illustrates the LBP, LTP and SLRB images of four images under different illumination conditions.

3. Experiments

3.1. Similarity measures

The SLRB descriptor of a face image is a bit string after binary encoding according to Eq. (11), thus we adopt the cosine similarity S_C and Hamming similarity S_H .

Suppose $\beta^{(1)} = [\beta_1^{(1)}, \beta_2^{(1)}, \dots, \beta_L^{(1)}]^T$ and $\beta^{(2)} = [\beta_1^{(2)}, \beta_2^{(2)}, \dots, \beta_L^{(2)}]^T$ are two SLRB strings with length L . The cosine similarity is defined as

$$S_C(\beta^{(1)}, \beta^{(2)}) = \frac{1}{\|\beta^{(1)}\| \|\beta^{(2)}\|} \sum_{l=1}^L \beta_l^{(1)} \beta_l^{(2)}, \quad (15)$$

which measures their similarity by their angles in an inner product space. However, it just reflects how often 1s coincide in each string. We employ another similarity using their Hamming distance D_H which can be formulated as

$$D_H(\beta^{(1)}, \beta^{(2)}) = \sum_{l=1}^L |\beta_l^{(1)} - \beta_l^{(2)}|. \quad (16)$$

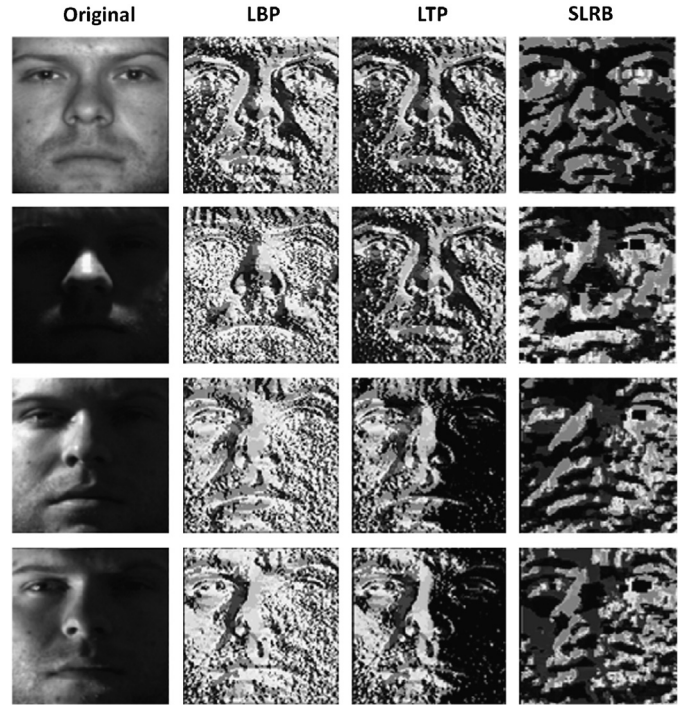


Fig. 5. Comparison of the LBP, LTP and SLRB images.

The maximum Hamming distance of two SLRB descriptors with length L is L . Therefore the Hamming similarity is defined as

$$S_H(\beta^{(1)}, \beta^{(2)}) = \left(1 - \frac{D_H(\beta^{(1)}, \beta^{(2)})}{L}\right)^2. \quad (17)$$

3.2. Experiment settings

Experiments are carried out on the Extended Yale Face Database B and CMU-PIE to illustrate the effectiveness of the SLRB descriptor. The Extended Yale Face Database B is an updated version of the Yale Face Database B, containing 38 subjects with 9 poses and 64 illumination conditions. The images are divided into five subsets according to the angle between the light source direction and the central camera axis. We utilized the images with the most neutral light sources as the gallery and ones from the five subsets as probes. The CMU-PIE face database includes 68 subjects with 41 368 images under variations in pose, illumination and expression. We chose the illumination subset (1425 images of 68 subjects under illumination from 21 directions) in our experiments. One image per subject were chosen as the gallery each time and the others were used as the probes. All face images were properly aligned, cropped and resized to 120×120 in the experiments.

Each image is divided into several blocks with the sizes 4×4 , 5×5 , 6×6 , 7×7 to obtain the SLRB features. The ℓ_1 -norm regularization parameter and the ℓ_2 -norm regularization parameter are set 1. We determine the linear regression coefficient vector α by three methods: the OLS estimation, ridge regression and SLRB (lasso regression), which are denoted as “OLS”, “RR” and “SLRB” respectively. We also overlap adjacent blocks to obtain more robust features which are denoted as “OLSo”, “RRo” and “SLRBo” respectively.¹ We adopt the cosine similarity and Hamming similarity which are denoted as “SLRBc” and “SLRBh” respectively. The SLRB

¹ In our experiments, the size of overlap areas is half of the blocks. For odd-size blocks, we round towards plus infinity.

Table 1

Recognition rates (%) on extended Yale Face Database B with cosine similarity.

cos	4 × 4	5 × 5	6 × 6	7 × 7
OLS	88.11	70.07	75.99	74.95
RR	87.48	89.24	85.95	83.42
SLRB	86.09	89.17	90.06	87.35
OLSo	89.92	79.07	83.38	80.88
RRo	88.84	90.11	88.21	85.43
SLRBo	88.44	90.58	91.93	89.15

Table 2

Recognition rates (%) on extended Yale Face Database B with Hamming similarity.

ham	4 × 4	5 × 5	6 × 6	7 × 7
OLS	88.61	70.46	76.87	75.11
RR	87.24	89.15	86.32	83.75
SLRB	86.66	87.76	88.37	85.06
OLSo	90.09	79.95	83.59	82.45
RRo	88.35	90.36	88.73	85.83
SLRBo	89.07	89.32	91.20	88.46

Table 3

Recognition rates (%) on extended Yale Face Database B.

Methods	Subset					Average
	S1	S2	S3	S4	S5	
LBP	100	100	96.92	61.03	34.87	78.56
LTP	100	100	97.80	76.62	58.40	86.56
SLRBc	100	100	91.42	83.65	84.59	91.93
SLRBh	100	100	87.03	83.26	85.71	91.20

methods are compared with the LBP and LTP methods. We use the nearest neighborhood rule with ℓ_2 -norm as the classifier.

3.3. Experimental results

Table 1 and Table 2 show recognition rates on the Extended Yale Face Database B with the cosine similarity and Hamming similarity respectively. SLRB (lasso regression) is superior to other two methods. The overlap method achieves better results. Table 3 shows the SLRB recognition performance on the five subsets of the Extended Yale Face Database B, the SLRB descriptor with cosine similarity achieves 13.37% and 5.37% higher than LBP and LTP respectively. The recognition rate of the SLRB method in the third subset containing face images with relatively good illumination is lower than those of LBP and LTP. It shows that our SLRB descriptor tends to perform better in removing the influence of the illumination changes with a modest ability for the controlled face recognition. It may be caused by the information loss when employing the binary encoding and the limited number of feature patterns when utilizing lasso regression. However in the last two subsets SLRB performs significantly better than LBP and LTP, which demonstrates its robustness to the severe illumination degradation. It is noticeable that our proposed methods perform comparably for subsets 4 and 5 (subset 5 even slightly better than subset 4). This may be explained in Fig. 6 by the fact that the cosine distance between (a) and (b) is a little bigger than the cosine distance between (a) and (c). Fig. 7 illustrates the rank-10 average recognition rate on the Extended Yale Face Database B. Our method achieves a more stable and satisfactory result than the other two.

The recognition rates of different methods on the CMU-PIE face database are illustrated in Table 4. The SLRB descriptor achieves the highest recognition rates which illustrates the advantage of our methods. Furthermore, we explore the performance of the proposed method when changing the face images with various illumination conditions as gallery images. Fig. 8 shows the recognition rates of different methods for each gallery on the CMU-PIE face database. The SLRB descriptor can significantly improve the recog-

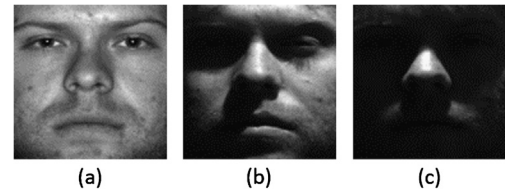


Fig. 6. (a) Gallery image; (b) a probe image in subset 4; (c) a probe image in subset 5. The cosine distance between (a) and (b) is 0.67 while the cosine distance between (a) and (c) is 0.65.

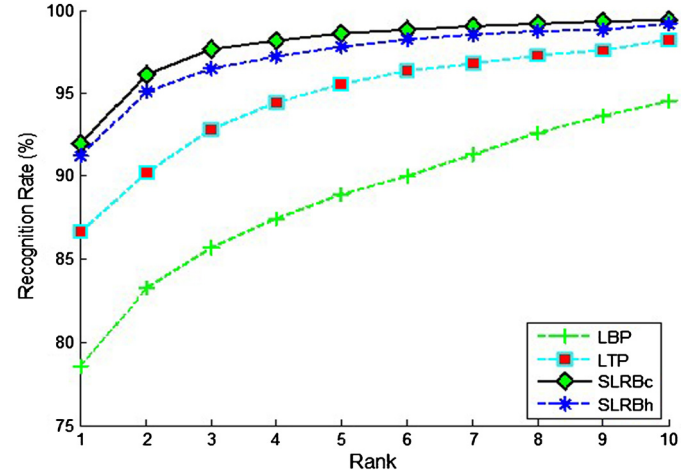


Fig. 7. The rank-10 recognition rate on the Extended Yale Database B using LBP, LTP and SLRB.

Table 4

Recognition rates (%) on CMU-PIE.

LBP	LTP	SLRBc	SLRBh
86.93	88.15	92.34	92.02

nition performance compared to LBP and LTP, especially when the gallery image is with severe illumination conditions.

Fig. 9 demonstrates the robustness to noise of different descriptors on the Extended Yale Face Database B. The white Gaussian noise with different variances (0.001, 0.002, 0.003, 0.004, 0.005) were added to the probe images. The recognition rate of SLRB drops about 20% while LTP drops 65% and LBP drops 75%. Our method shows more robustness to the noise due to the result of sparse regression and binary encoding.

One of our SLRB advantages is computational efficiency so we compare the CPU time on the Extended Yale Face Database B of SLRB with other algorithms. For an 8-bit gray-scale image with the size 120×120 , the features of LBP and LTP are 13 275 bits and 115 200 bits separately while the features based on the SLRB and SLRBo descriptors are 3200 bits and 12 800 bits separately when the size of blocks is 6×6 . The results are illustrated in Table 5 from which we could conclude that our method is more computationally efficient than other algorithms since the SLRB descriptor is a bit string with a low dimension and the SLRBo descriptor sacrifices computational efficiency for a higher recognition rate.

4. Conclusion

The SLRB descriptor has been proposed based on the assumption of locally linear consistency and Lambertian reflectance model in this paper. It is an illumination-insensitive representation and shows more robust to noise on account of the sparse regression and binary encoding. Experimental results on Extended Yale-B Database and CMU-PIE Database have shown that our approach

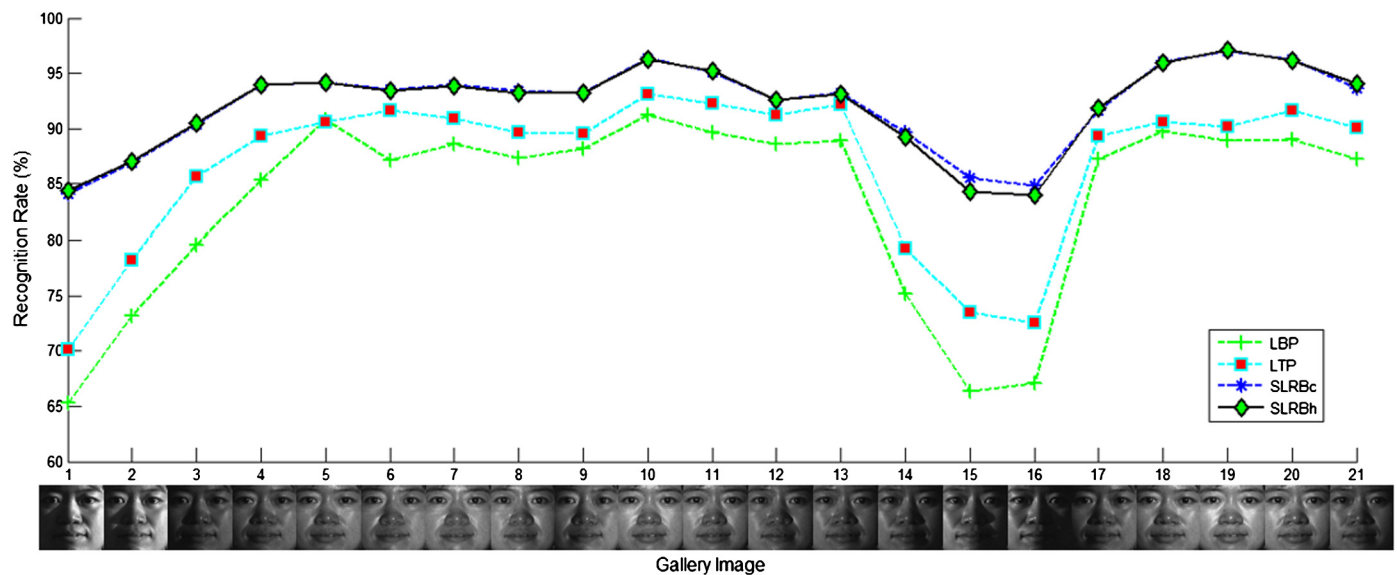


Fig. 8. Recognition rates versus different gallery images with various illumination conditions.

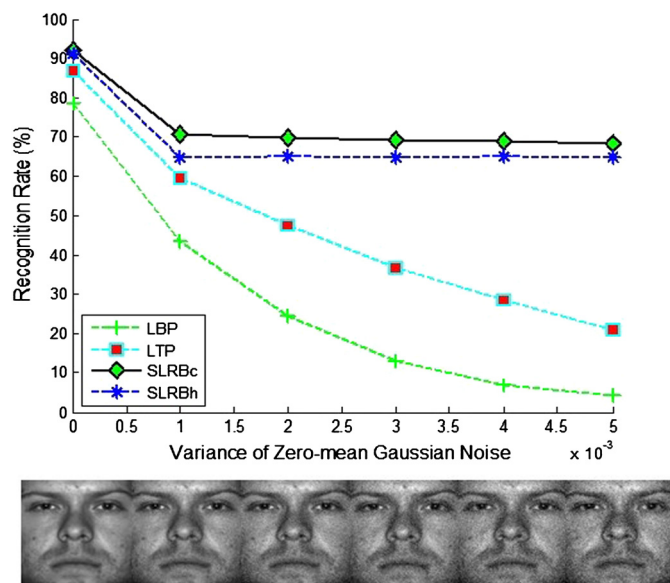


Fig. 9. Illustration of the robustness to illumination variations with different methods on the Extended Yale Database B.

Table 5
Average CPU time (in ms) per image on extended Yale Face Database B.

LBP	LTP	SLRB	SLRBc
15.0	13.4	6.8	26.8

yields better performance than several other approaches. However it has some drawbacks. Because of the nature of the bit-string arrangement, our method is not invariant to spatial transformation and requires a precise registration process. The SLRB descriptor is designed for bad illumination conditions and has a modest ability for the controlled face recognition. We will consider these as the direction of our future work.

References

[1] W. Zhao, R. Chellappa, P. Phillips, A. Rosenfeld, Face recognition: a literature survey, *ACM Comput. Surv.* 35 (4) (2003) 399–458.

[2] Y. Adini, Y. Moses, S. Ullman, Face recognition: the problem of compensating for changes in illumination direction, *IEEE Trans. Pattern Anal. Mach. Intell.* 19 (7) (1997) 721–732.

[3] X. Zou, J. Kittler, K. Messer, Illumination invariant face recognition: a survey, in: *IEEE International Conference on Biometrics: Theory, Applications, and Systems*, BTAS 2007, IEEE, 2007, pp. 1–8.

[4] S. Pizer, E. Amburn, J. Austin, R. Cromartie, A. Geselowitz, T. Greer, B. ter Haar Romeny, J. Zimmerman, K. Zuiderveld, Adaptive histogram equalization and its variations, *Comput. Vis. Graph. Image Process.* 39 (3) (1987) 355–368.

[5] S. Shan, W. Gao, B. Cao, D. Zhao, Illumination normalization for robust face recognition against varying lighting conditions, in: *IEEE International Workshop on Analysis and Modeling of Faces and Gestures*, AMFG 2003, IEEE, 2003, pp. 157–164.

[6] M. Savvides, B. Kumar, Illumination normalization using logarithm transforms for face authentication, in: *Audio- and Video-Based Biometric Person Authentication*, Springer, 2003, pp. 549–556.

[7] P. Hallinan, A low-dimensional representation of human faces for arbitrary lighting conditions, in: *Proceedings of IEEE Computer Society Conference on Computer Vision and Pattern Recognition*, CVPR'94, IEEE, 1994, pp. 995–999.

[8] A. Georgiades, P. Belhumeur, D. Kriegman, From few to many: illumination cone models for face recognition under variable lighting and pose, *IEEE Trans. Pattern Anal. Mach. Intell.* 23 (6) (2001) 643–660.

[9] R. Basri, D. Jacobs, Lambertian reflectance and linear subspaces, *IEEE Trans. Pattern Anal. Mach. Intell.* 25 (2) (2003) 218–233.

[10] M. De Marsico, M. Nappi, D. Riccio, H. Wechsler, Robust face recognition for uncontrolled pose and illumination changes, *IEEE Trans. Syst. Man Cybern., Part A, Syst. Hum.* 43 (1) (2013) 149–163.

[11] D. Jobson, Z.-U. Rahman, G. Woodell, Properties and performance of a center/surround retinex, *IEEE Trans. Image Process.* 6 (3) (1997) 451–462.

[12] D. Jobson, Z.-U. Rahman, G. Woodell, A multiscale retinex for bridging the gap between color images and the human observation of scenes, *IEEE Trans. Image Process.* 6 (7) (1997) 965–976.

[13] A. Baradarani, Q.J. Wu, M. Ahmadi, An efficient illumination invariant face recognition framework via illumination enhancement and dd-dtcwt filtering, *Pattern Recognit.* 46 (1) (2013) 57–72.

[14] R.C. Gonzales, R.E. Woods, *Digital Image Processing*, 2nd ed., Prentice-Hall, Upper Saddle River, NJ, 2002.

[15] W. Chen, M. Er, S. Wu, Illumination compensation and normalization for robust face recognition using discrete cosine transform in logarithm domain, *IEEE Trans. Syst. Man Cybern., Part B, Cybern.* 36 (2) (2006) 458–466.

[16] S. Du, R. Ward, Wavelet-based illumination normalization for face recognition, in: *IEEE International Conference on Image Processing*, vol. 2, ICIP 2005, IEEE, 2005, pp. II–954.

[17] H. Wang, S. Li, Y. Wang, Face recognition under varying lighting conditions using self quotient image, in: *IEEE International Conference on Automatic Face and Gesture Recognition*, IEEE, 2004, pp. 819–824.

[18] S. Wei, S. Lai, Robust face recognition under lighting variations, in: *Proceedings of the 17th International Conference on Pattern Recognition*, vol. 1, ICPR 2004, IEEE, 2004, pp. 354–357.

[19] T. Zhang, Y. Tang, B. Fang, Z. Shang, X. Liu, Face recognition under varying illumination using gradientfaces, *IEEE Trans. Image Process.* 18 (11) (2009) 2599–2606.

- [20] B. Wang, W. Li, W. Yang, Q. Liao, Illumination normalization based on Weber's law with application to face recognition, *IEEE Signal Process. Lett.* 18 (8) (2011) 462–465.
- [21] Y. Wu, Y. Jiang, Y. Zhou, W. Li, Z. Lu, Q. Liao, Generalized Weber-face for illumination-robust face recognition, *Neurocomputing* 136 (2014) 262–267.
- [22] T. Ojala, M. Pietikainen, T. Maenpää, Multiresolution gray-scale and rotation invariant texture classification with local binary patterns, *IEEE Trans. Pattern Anal. Mach. Intell.* 24 (7) (2002) 971–987.
- [23] X. Tan, B. Triggs, Enhanced local texture feature sets for face recognition under difficult lighting conditions, *IEEE Trans. Image Process.* 19 (6) (2010) 1635–1650.
- [24] S. Roweis, L. Saul, Nonlinear dimensionality reduction by locally linear embedding, *Science* 290 (5500) (2000) 2323–2326.
- [25] R. Tibshirani, Regression shrinkage and selection via the lasso, *J. R. Stat. Soc. B* (1996) 267–288.

Zuodong Yang received the B.E. degree from Department of Electronics and Information Engineering, Tsinghua University, Beijing, China, in 2013. He is currently working towards the Master degree in electronics engineering at Tsinghua University, Beijing, China. His research interests include image processing and pattern recognition in biometrics.

Yong Wu received the B.E. degree from Department of Electronics and Information Engineering, Tsinghua University, Beijing, China, in 2011. He is currently working towards the Master degree in electronics engineering at Tsinghua University, Beijing, China. His research interests include image processing and pattern recognition in biometrics.

Wenteng Zhao received the B.E. degree from Department of Electronics and Information Engineering, Tsinghua University, Beijing, China, in 2013. He is currently working towards the Master degree in electronics engineering at Tsinghua University, Beijing, China. His research interests include image processing.

Yicong Zhou received his B.S. degree from Hunan University, Changsha, China, and his M.S. and Ph.D. degrees from Tufts University, Massachusetts, USA, all degrees in electrical engineering. He is currently an Assistant Professor in the Department of Computer and Information Science at Uni-

versity of Macau, Macau, China. His research interests focus on multimedia security, image/signal processing, pattern recognition and medical imaging. Dr. Zhou is a member of the IEEE and SPIE (International Society for Photo-Optical Instrumentations Engineers).

Zongqing Lu received the Ph.D. degree in signal processing from XiDian University Xi'an China in 2007. From 2007 to 2010, he was a post-doctoral fellow at Tsinghua University. Now he is an assistant professor in the Electronic engineering Department of Tsinghua University. His research interests include image processing, machine learning.

Weifeng Li received the M.E. and Ph.D. degrees in Information Electronics at Nagoya University, Japan, in 2003 and 2006, respectively. He joined the Idiap Research Institute, Switzerland, in 2006, and in 2008 he moved to Swiss Federal Institute of Technology, Lausanne (EPFL), Switzerland, as a research scientist. Since 2010 he has been an associate professor in the Department of Electronic Engineering/Graduate School at Shenzhen, Tsinghua University, China. His research interests lie in the areas of audio and visual signal processing, Biometrics, Human–Computer Interactions (HCI), and machine learning techniques. He is a member of the IEEE and IEICE.

Qingmin Liao received the B.S. degree in radio technology from the University of Electronic Science and Technology of China, Chengdu, China, in 1984, and the M.S. and Ph.D. degrees in signal processing and telecommunications from the University of Rennes 1, Rennes, France, in 1990 and 1994, respectively. Since 1995, he has been joining with Tsinghua University, Beijing, China. In 2002, he became Professor in the Department of Electronic Engineering of Tsinghua University. Since 2010, he has been the director of the Division of Information Science and Technology in the Graduate School at Shenzhen, Tsinghua University. He is also affiliated with the Shenzhen Key Laboratory of Information Science and Technology (Director), China. Over the last 30 years, he has published over 100 peer-reviewed journal and conference papers. His research interests include image/video processing, transmission and analysis; biometrics; and their applications to teledetection, medicine, industry, and sports.