

# Recognizing Partially Occluded, Expression Variant Faces From Single Training Image per Person With SOM and Soft $k$ -NN Ensemble

Xiaoyang Tan, Songcan Chen, Zhi-Hua Zhou, *Member, IEEE*, and Fuyan Zhang

**Abstract**—Most classical template-based frontal face recognition techniques assume that multiple images per person are available for training, while in many real-world applications only one training image per person is available and the test images may be partially occluded or may vary in expressions. This paper addresses those problems by extending a previous local probabilistic approach presented by Martinez, using the self-organizing map (SOM) instead of a mixture of Gaussians to learn the subspace that represented each individual. Based on the localization of the training images, two strategies of learning the SOM topological space are proposed, namely to train a single SOM map for all the samples and to train a separate SOM map for each class, respectively. A soft  $k$  nearest neighbor (soft  $k$ -NN) ensemble method, which can effectively exploit the outputs of the SOM topological space, is also proposed to identify the unlabeled subjects. Experiments show that the proposed method exhibits high robust performance against the partial occlusions and variant expressions.

**Index Terms**—Face expression, face recognition, occlusion, self-organizing map (SOM), single training image per person.

## I. INTRODUCTION

AS ONE of the few biometric methods that possess the merits of both high accuracy and low intrusiveness, face recognition technology (FRT) has a variety of potential applications in information security, law enforcement and surveillance, smart cards, access control, among others [1]–[3]. For this reason, FRT has received significantly increased attention from both the academic and industrial communities during the past twenty years. Numerous recognition methods have been proposed, some of which have obtained much success under constrained conditions [3]. However, the general face recognition problem is still unsolved due to its inherent complexity.

Manuscript received January 12, 2004; revised November 1, 2004. The work was supported in part by the National Natural Science Foundation of China under Grant 60271017, in part by the National Outstanding Youth Foundation of China under Grant 60325207, in part by the Jiangsu Science Foundation Key Project under Grant BK2004001, in part by the Ying Tung Education Foundation under Grant 91067, and in part by the Excellent Young Teachers Program of MOE of China.

X. Tan is with the National Laboratory for Novel Software Technology, Nanjing University, Nanjing 210093, China. He is also with the Department of Computer Science and Engineering, Nanjing University of Aeronautics and Astronautics, Nanjing 210016, China (e-mail: x.tan@nuaa.edu.cn).

S. Chen is with the Department of Computer Science and Engineering, Nanjing University of Aeronautics and Astronautics, Nanjing 210016, China. He is also with the Shanghai Key Laboratory of Intelligent Information Processing, Fudan University, Shanghai 200433, China (e-mail: s.chen@nuaa.edu.cn).

Z.-H. Zhou and F. Zhang are with the National Laboratory for Novel Software Technology, Nanjing University, Nanjing 210093, China (e-mail: zhouzh@nju.edu.cn, fy.zhang@nju.edu.cn).

Digital Object Identifier 10.1109/TNN.2005.849817

The complexities of face recognition mainly lie in the constantly changing appearance of human face, such as variations in occlusion, illumination and expression. One way to overcome these difficulties is to explain the variations by explicitly modeling them as free parameters [4], [5]. On the other hand, the variations can be attacked indirectly by searching one or more face subspaces of the face so as to lower the influence of the variations. The early way to construct face subspaces is by manually measuring the geometric configurational features of the face [8]. Later, researchers realize that one can extract not only the texture and shape information but also the configurational information of the face directly from its raw-pixels-based representation [9]. In addition, a higher recognition rate can also be achieved with latter method as compared to the geometric-based approach [8]. Consequently, template-based methods have become one of the dominant techniques in the field of face recognition since the 1990s.

To construct the face subspaces with good generalization, a large and representative training data set should be required due to the high dimensions of the face images [10]. However, this is not always possible in many real world tasks, such as finding a person within a large database of faces (e.g., in the law enforcement scenarios), typically only one image is available per person. This brings much trouble to many existing algorithms. Most subspace methods such as linear discriminant analysis (LDA) [12]–[14], Bayesian matching methods [16] and evolutionary pursuit (EP) [35] may fail because the intra-class variation becomes zero when there is only one image per class for training. Furthermore, even the parameter estimation becomes possible when two or more samples are available, these algorithms' performance may still suffer a lot from the small, non-representative sample sets [17], [18].

In fact, this so-called *one image per person* problem can be traced back to the early period when the geometric-based methods were popular, where various configurational features such as the distance between two eyes are manually extracted from the single face image [8]. Recently, several researchers [11], [18]–[20] have begun to again pay attention to this classical problem within the template-based framework due to the needs of the applications and its potential significance of solving the likewise small sample problem.

The general idea behind these literatures to solve this problem is to try to squeeze as much information as possible from the single face image, which is used to provide each person with several imitated face images. For example, Chen *et al.* enlarged the training image database using a series of  $n$ -order projected

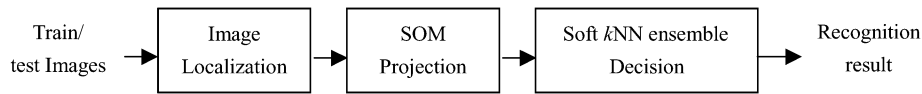


Fig. 1. Block diagram of the proposed method for face recognition.

images [19]. Beymer and Poggio developed a method to generate virtual views by exploiting prior knowledge of faces to deal with the pose-invariant problem [20]. Improved recognition accuracies have been achieved by these methods. However, as Martinez stated [18], one nonignorable drawback of these methods is that it may be highly correlated among the generated virtual images and therefore these samples should not be considered as independent training images.

Besides the one image per person problem, there exist other problems that make things even more complicated, such as the *partial occlusion* and/or *expression-invariant* problem. The former means that portions of the face image are missed or occluded due to glasses or clothing, while the latter means that the facial expressions in the training images are different from those in the testing images of the same subject. Both the problems are supposed to be difficult for the traditional template-based paradigm, such as principal component analysis (PCA) [27]–[29].

Recently, Martinez [18], [21], [22] has partially tackled the previous problems using a local probabilistic method, where the subspace of each individual is learned and represented by a separate Gaussian distribution. This paper extends his work by proposing an alternative way of representing the face subspace with self-organizing maps (SOMs) [24], [25]. One of the main motivations of such an extension is that, even when the sample size is too small to faithfully represent the underlying distribution (e.g., when not enough or even no virtual samples are generated), the SOM algorithm can still extract all the significant information of local facial features due to the algorithm's unsupervised and nonparametric characteristic, while eliminating possible faults like noise, outliers, or missing values. In this way, the compact and robust representation of the subspace can be reliably learned. Furthermore, a soft  $k$  nearest neighbor (soft  $k$ -NN) ensemble method, which can efficiently exploit the outputs of the SOM topological space, is also proposed to identify the unlabeled subjects.

The paper proceeds as follows. The local probabilistic method proposed by Martinez is briefly introduced in Section II. The proposed method is described in Section III. The experiments are reported in Section IV. Finally, the conclusion is drawn in Section V.

## II. LOCAL PROBABILISTIC APPROACH

The local probabilistic approach works as follows [18], [21]. First, a set of virtual images accounting for all possible localization errors of the original image are synthetically generated for each training face. And then, each face (including the generated face images) is divided into six local areas, and the subspace of every subimage is estimated by means of a mixture model of Gaussians using the EM algorithm. Finally, the eigenrepresentation of each local areas are calculated within their own subspace, and each sample image can thus be represented as a mixture of Gaussians in each of these eigenspaces.

In the identification stage, the test images are also divided into six local areas and are then projected onto the above computed

eigenspaces. A probabilistic rather than a voting approach is used to measure the similarity of a given match. Experiments on a set of 2600 images show that the local probabilistic approach does not reduce accuracy when 1/6 of the face is occluded (on the average).

A weighted local probabilistic method is also proposed by the author to address the expression-invariant problem [18], [22]. The idea is based on the fact that different facial expressions influence different parts of the face more than others, thus, a learning mechanism is proposed to learn such prior information by weighting each of the local parts with the values that vary depending on the facial expression displayed on the testing image.

The (weighted) local probabilistic method greatly improves the robustness of the recognition system. However, the mixture of Gaussians is a parametric method which heavily depends on the assumption that the underlying distribution is faithfully represented with a lot of samples. While those samples can be synthetically generated as the way described previously, the computational and storage costs along with the procedure of generating virtual samples may be very high (e.g., 6,615 samples per individual in [18]) when the face database is very large. We extend the method using an unsupervised and nonparametric method, i.e., SOM, which can represent the subspace of local features reliably even no extra virtual samples are generated, and the application scope of the method is hereby expanded. The proposed method will be detailed in Section III.

## III. PROPOSED METHOD

A high-level block diagram of the proposed method is shown in Fig. 1. The details of the method are described in Section III-A–C.

### A. Localizing the Face Image

In the case of only limited training samples available per person, it is almost unavoidable to face the dilemma of high dimensions of image data and small samples. One way to deal with this difficulty is to reduce the dimensionality using projection methods such as PCA [27]–[29], however, one of the main drawbacks of the classical PCA projection is that it is easy to be subject to gross variations and, thus, sensitive to any changes in expression, illumination, etc. Most recently, Yang *et al.* have proposed a new technique named two-dimensional principal component analysis (2-DPCA) [43], which is based on small 2-D image matrices instead of 1-D vectors and, thus, more suitable for small sample size problem than classical PCA.

An alternative way is to use local approaches [8], [11], [16], [23], [28], [30], [31], [34], [48], i.e., based on some partition of the image, the original face can be represented by several *low* dimensional local feature vectors (LFVs) rather than one full *high* dimensional vector, thus, the small sample problem can be alleviated. Moreover, the subpattern dividing process can also help

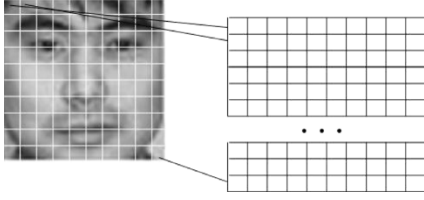


Fig. 2. Representation of a face image by a set of subblock vectors.

increase the diversity [31], making the common and class-specific local features be identified more easily [30]. These advantages are useful for the identification of the unknown persons.

In the implementation of this paper, the original image is divided into  $M (= l/d)$  nonoverlapping subblocks with equal size. The  $M$  LFVs are obtained by concatenating the pixels of each subblock, where  $l$  and  $d$  are the dimensionalities of the whole image and each subblock, respectively (Fig. 2). Note that other local features can also be used, such as the local eigenfeatures [16], [21], [28] and local Gabor filters [36]. For example, one can extract local features at different scale as done in [50] to deal with the multiscale problem.

### B. Use of SOM

1) *Single SOM-Face Strategy*: A large number of subblocks may be generated from the previous localizing stage. Now we need an efficient method to find the patterns or structures of the subblocks without any assumption about their distribution. There exist several classical methods for this purpose, such as the LGB [32], SOM [24], [25], and fuzzy C-means [33]. Fleming *et al.*'s early work also suggested that unsupervised feature clustering helps to improve the performance of the recognition system [37].

The SOM, which approximates an unlimited number of input items by a finite set of weight vectors, is chosen here for several reasons as follows. First, the SOM learning is efficient and effective, suitable for high-dimensional process [25]. Second, it is found that the SOM algorithm is more robust to initialization than other algorithm such as LGB [38]. Finally and most importantly, in addition to clustering, the weight vectors can be organized in an ordered manner so that the topologically close neurons are sensitive to similar input subblocks [25].

In this paper, the original SOM algorithm is used, which is arguably one of the most computationally efficient [39]. This is very important, especially when the number of subblocks is very large. Furthermore, to accelerate the computation of the SOM, the batch formulation of the SOM algorithm is used, which has the advantages of fast convergence and being less sensitive to the order of presentation of the data. Recently, the batch-SOM algorithm has been employed by Kohonen *et al.* to fast learn a large text corpus in their WEBSOM project [39].

Formally, let  $X(t) = \{x_i(t) | i = 1, \dots, N\}$  be the set of input vectors at time  $t$ , and  $A = \{e_1, e_2, \dots, e_Q\}$  be the neurons of the SOM, respectively. Also let the weight vector (also called codebook or reference vectors) stored in the neuron  $e_i (i \in [1, Q])$  be  $w_i$ , which in turn decides the location of the neuron  $e_i$  in the lattice  $A$ . The Voronoi set of the neuron  $e_i$  is denoted by  $V_i$ , which consists of all the  $n_i$  closest subblocks of the weight vector  $w_i$  in the input space. Let  $N$  denote the size of the data set and  $Q$ ,

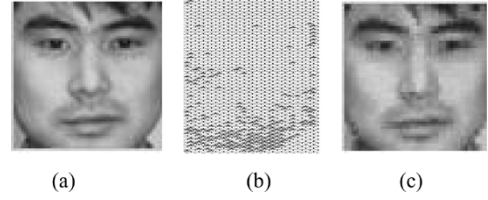


Fig. 3. Example of an (a) original image, (b) its projection, and (c) the reconstructed image.

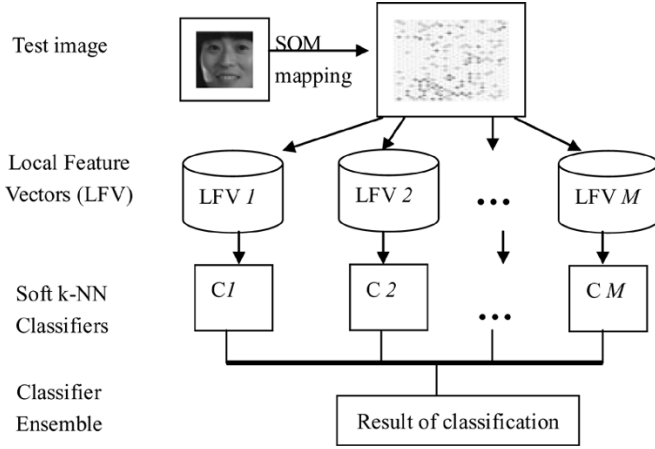
the number of neurons in the lattice, respectively. After initialization, the batch-mode SOM algorithm consists of iterating the following three steps to adjust the weight vectors until they can be regarded as stationary [25], [40].

- Step 1) Partition all the subblocks into Voronoi regions  $V_i (i \in [1, Q])$  by finding each subblock's closest weight vector according to:  $c = \arg \min_i \{\|x_i(t) - w_i(t)\|\}$ , where  $c$  is index of the winner neuron.
- Step 2) And then the average of the subblock vectors  $x(t)$  over  $V_i$ , denoted by  $\bar{x}_i$ , is computed by:  $\forall i, \bar{x}_i = \sum_{x(t) \in V_i} x(t) / n_i$ , with  $N = \sum_i n_i$
- Step 3) Smooth the weight vector of each neuron:  $w_i^* = \sum_j n_j h_{jr} \bar{x}_j / \sum_j n_j h_{jr}$ , where  $h_{j,r}$  is a neighborhood function which governs both the self-organization process and the topographic properties of the map.

After the SOM map has been trained, all the subblocks from each training face are mapped to the best matching units (BMUs) in the SOM topological space by a nearest neighbor strategy. The corresponding weight vectors of the BMUs will be used as the "prototype" vectors of each class for later recognition purpose (detailed in Section III-B.2). Fig. 3 shows an example of an original image, its projection and the reconstructed image (called "SOM-face," constructed with the corresponding prototype vectors).

2) *Multiple SOM-Face Strategy*: One drawback of the manifold-based learning methods is the requirement of the recomputation of the base vectors (such as eigenvectors in PCA, weight vectors in SOM, etc.) when new individuals are presented. One way to deal with this problem is to train a separate subspace for each face, that is, when a new face is encountered, only the new one rather than the original whole database is needed to be relearned, thus, the computational cost can be significantly reduced due to the simplification of the learning procedure.

Based on the previous idea, another SOM-face strategy, namely the multiple SOM-face (MSOM-face) strategy is proposed (similar ideas can be found in [15] and [26]). Formally, let the  $j$ th new face image be presented to the system, which is divided into a set of equally-sized subblocks, denoted as  $\{x_i^j | i = 1, \dots, M\}$ . Then a separate small SOM map  $S^j$  for the face will be trained with its class-designated samples, using the previously-described batch-mode SOM algorithm. To evaluate the prototype vectors of the new class, we can map the new face's subblocks onto the newly trained map. Another way is to combine all the small maps into one big map and then recalculate all the prototype vectors for every class. We have found that the latter method is more accurate when a large number of classes exist.

Fig. 4. Architecture of the soft  $k$ -NN ensemble classifier.

### C. Identify Faces Based on SOM-Face

To identify an unlabeled face, a classifier should be built on the SOM map. Popular supervised classifiers, such as LVQ, MLP, SVM, etc, require that the neurons of the SOM be labeled before being used for recognition. Unfortunately, labeling neuron may also lead to the losing of information. For example, in a commonly used labeling mechanism, majority voting (MV), each neuron is hard labeled according to the maximum class frequency obtained from the training data, while a large amount of information about the classes other than the winner class is, thus, lost from the neuron.

For this reason, a soft  $k$ -NN ensemble decision scheme (Fig. 4) is proposed to avoid the above problem and to effectively exploit as much information as possible from the outputs of the SOM topological space. In the proposed method, each component classifier is supposed to output a confidence vector which gives the degree of support for each LFVs membership in every class, and then all the outputs will be combined with a sum aggregation method [49] to give the final decision. The details of the proposed method are described in the following.

Given  $C$  classes, to decide which class the test face  $x$  belongs to, we first divide the test face into  $M$  nonoverlapping sub-blocks as mentioned before, and then project those subblocks onto the trained SOM maps to obtain the face's SOM-face representations.

Then, a distance matrix, describing the dissimilarity between the test face and every training face in the current SOM topological space, is calculated. The distance-calculation algorithm will be detailed in Table I. Its outputs, however, are introduced directly below for convenience of description, which take the form of a  $M \times C$  matrix as follows:

$$D(x) = \begin{bmatrix} d_{11} & d_{12} & \dots & d_{1C} \\ d_{21} & d_{22} & \dots & d_{2C} \\ \dots & \dots & \dots & \dots \\ d_{M1} & d_{M2} & \dots & d_{MC} \end{bmatrix} \quad (1)$$

$$\equiv [dv_1 \dots dv_M]^T \quad (2)$$

$$dv_j = [d_{j1}, d_{j2}, \dots, d_{jC}]$$

where  $dv_j$  is called the distance vector, whose element  $d_{jk}$  is the distance between the  $j$ th neuron of the test face and the corresponding neuron of the  $k$ th class.

TABLE I  
DISTANCE MATRIX CALCULATION

|                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                        |
|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| Step1: Given $C$ training images, the $i$ -th image's BMU (Best Matching Unit) set is denoted by $Train_i = \{n_1^i, n_2^i, \dots, n_M^i\}$ , where $i = 1 \dots C$ , and $M$ is the number of sub-blocks of each face.                                                                                                                                                                                                                                                                                                                                                                                                                                                                |
| Step2: Divide the test image $x$ into $M$ non-overlapping sub-blocks with equal size.                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                  |
| Step3: 1) For Single SOM-face strategy: <ul style="list-style-type: none"> <li>Project <math>M</math> sub-blocks of the test image onto the single trained map, the obtained BMU set is denoted by <math>Test_s = \{n_1^s, n_2^s, \dots, n_M^s\}</math>.</li> <li>For <math>i = 1 \dots M</math> <ol style="list-style-type: none"> <li>(1) Calculate the distance between <math>n_i^s</math> in <math>Test_s</math> and its counterparts <math>n_i^k</math> in <math>Train_k, k = 1 \dots C</math>.</li> <li>(2) The resulting distances are denoted by <math>\{d_{i,k}, k = 1 \dots C\}</math>.</li> </ol> </li> </ul>                                                               |
| 2) For MSOM-face strategy: <ul style="list-style-type: none"> <li>Project <math>M</math> sub-blocks of the test image onto the <math>C</math> trained maps separately, denote each of the obtained BMU set by <math>Test_s^k = \{n_1^k, n_2^k, \dots, n_M^k\}, k = 1 \dots C</math>.</li> <li>For <math>k = 1 \dots C</math>: <ol style="list-style-type: none"> <li>(1) In <math>k</math>-th map, calculate the distance between each <math>n_i^k</math> in <math>Test_s^k</math> and its counterparts <math>n_i^i</math> in <math>Train_i, i = 1 \dots M</math>.</li> <li>(2) The resulting distances are denoted by <math>\{d_{i,k}, i = 1 \dots M\}</math>.</li> </ol> </li> </ul> |
| Step4: arrange the obtained distances into an $M \times C$ distance matrix $D$ , with its element being $\{d_{ik}, k = 1 \dots C, i = 1 \dots M\}$ .                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                   |

Next, we convert the distance vectors into corresponding confidence vectors. One possible conversion method is to use a soft  $k$ -NN algorithm, which is briefly described as follows. The distance from the  $j$ th neuron of the test face to its  $k$ -NNs are first arranged in increasing order:  $d_{j1}^* \leq d_{j2}^* \leq \dots \leq d_{jk}^*$ , then the confidence value for the  $k$ th nearest neighbor is defined as

$$c_{jk} = \frac{\log(d_{j1}^* + 1)}{\log(d_{jk}^* + 1)}. \quad (3)$$

It can be seen from (3) that the class with minimum distance to the testing subblock will yield a confidence value closer to one, while a large distance produces a very small confidence value, meaning that it is less likely for the testing subblock to belong to that class.

Finally, the label of the test image can be obtained through a linearly weighted voting scheme, as follows:

$$\text{Label} = \arg \max_k \left( \sum_{j=1}^M c_{jk} \right), \quad k = 1 \dots C. \quad (4)$$

## IV. EXPERIMENTS

### A. Data Set and Methodology

To verify the performance of the proposed method, we have conducted various experiments on two well-known face image databases (AR [44], [45] and FERET [46]). The AR database is employed to test the performance of the SOM-face algorithm under the conditions when partial occlusion and expression variant are involved. The FERET database is used to explore some practical properties of the proposed algorithm, such as the selection of the size of subblock.

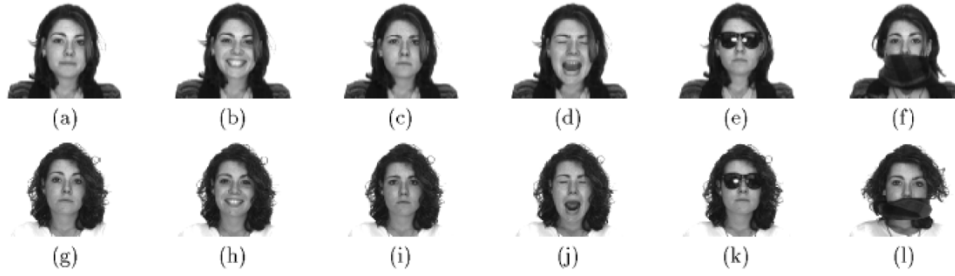


Fig. 5. Sample images for one subject of the AR database.

Except when stated otherwise, all the experiments reported here were conducted as follows. In the localizing step, the training images were partitioned into nonoverlapping subblock with equal size. Then either a single SOM map or multiple SOM maps was trained in batch-mode using the obtained subblocks. The training process was divided into two phases as recommended by [25], that is, a rough training phase (to adjust the topological order of the weight vectors) and a fine-adjustment phase (to fine tune the feature map so as to provide an accurate statistical quantification of the input space). 100 updates were performed in the first phase, while 400 times in the second one. The initial weights of all neurons were set to the greatest eigenvectors of the training data, and the neighborhood widths of the neurons converged exponentially to 1 with the increase of training time. The soft  $k$ -NN ensemble classifier described before was used for final classification decision.

Finally, to evaluation the experimental results of the SOM-face algorithm compared to other methods, we have conducted pairwise one-tail statistical test using the method described in [47]. Recently, several other authors have also adopted the same statistical testing method [43].

### B. Experiments on the AR Database

The AR face database [44], [45] contains over 4000 color face images of 126 people's faces (70 men and 56 women), including frontal view faces with different facial expressions, illumination conditions, and occlusions (sun glasses and scarf). There are 26 different images per person, taken in two sessions (separated by two weeks), each session consisting of 13 images. In our experiments, a subset of 1200 images from 100 different subjects (50 males and 50 females) were used, which were the same dataset used by Martinez *et al.* in their experiments [18], [21]. All the experimental results reported in this section concerning the local probabilistic approach were also quoted directly from [18] and [21]. Some sample images for one subject are shown in Fig. 5.

Before the recognition process, each image was cropped and resized to  $120 \times 165$  pixels and then converted to gray-level images, which were then processed by a histogram equalization algorithm. Except when stated otherwise, in our experiments, the subblock size of  $5 \times 3$  was used and only the single SOM strategy was employed. According to [46], both the top 1 match and the top  $k$  matches were considered through the testing procedure.

1) *Variations in Facial Expressions:* First, we conducted experiments to investigate the classification capability of the proposed method under varying facial expressions. The results were

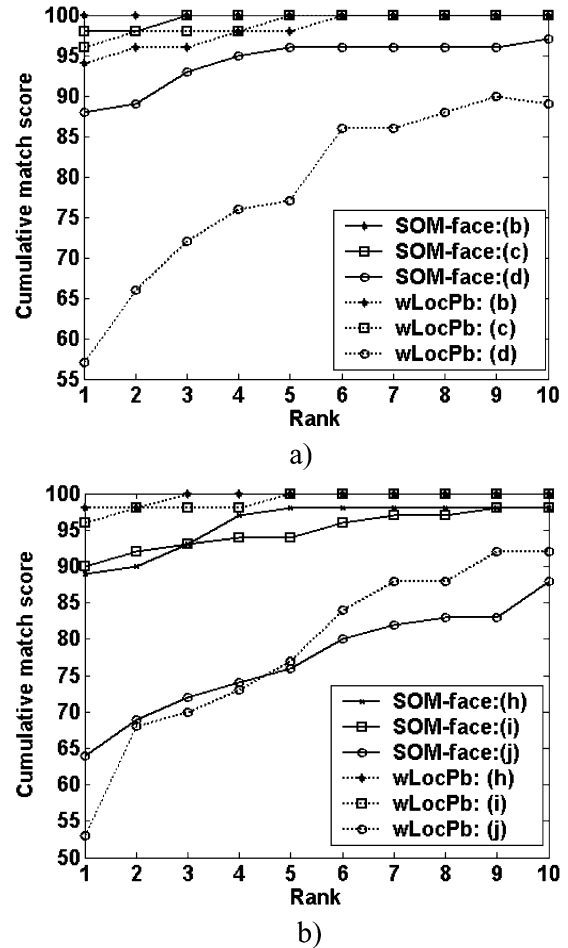


Fig. 6. Comparative performance of the SOM-face algorithm and the weight local probabilistic approach (wLocPb) [18] with expression variant involved.

compared with those of the weighted local probabilistic approach [18]. Specifically, in this series of experiments, the neutral expressions images [Fig. 5(a)] of the 100 individuals were used for training, while the smile, anger and scream images in the first sessions [Fig. 5(b)–(d)] and those in the second sessions [Fig. 5(h)–(j)] were used for testing. Thus, there were 300 testing images in all for each experiment.

Results of the two experiments are shown in Fig. 6(a) and (b), respectively, where the horizontal axis is rank and the vertical axis is cumulative match score, representing the percentage of correct matches with the correct answer in the top  $k$  matches. It can be seen from Fig. 6 that, when the happy and angry images [Fig. 5(b), (c), (h), and (i)] were presented to the system, the performance of the two methods is comparable. However, when

TABLE II  
COMPARISON OF PCA, 2-DPCA, AND SOM-FACE ALGORITHM CONCERNING  
THE TOP 1 RECOGNITION ACCURACY (%)

| Index | PCA | 2DPCA | SOM-face |
|-------|-----|-------|----------|
| (b)   | 97  | 99    | 100      |
| (c)   | 91  | 94    | 98       |
| (d)   | 63* | 69*   | 88       |
| (h)   | 67* | 72*   | 88       |
| (i)   | 63* | 69*   | 90       |
| (j)   | 40* | 46*   | 64       |

The asterisks indicate a statistically significant difference between SOM-face and the compared method at a significance level of 0.05 in the trials.

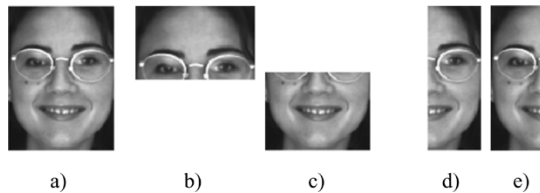


Fig. 7. Illustration of the image in different occlusion modes. Image (a) is the testing image without occlusion, while images (b)–(e) are faces in four occlusion modes, respectively.

the “screaming” faces [Fig. 5(d) and (j)], which differ most from the training sample, were used for testing, the SOM-face algorithm achieved much higher performance than the weighted local probabilistic approach, to be more specific, 30% and 10% higher accuracy is yielded, respectively, concerning the *top 1 match rate*. Moreover, it should be noted that the weighted local probabilistic approach must use additional training samples to obtain the needed weighted information of each local area for each expression, while the SOM-face achieves comparable or better recognition performance by using only one reference facial image per individual. Therefore, the proposed method seems to have a higher degree of robustness to variant facial expressions.

For reference, we present the top 1 recognition rates one would obtain with the classical PCA and 2-DPCA [43] in Table II.

2) *Variations in Partially Occluded Conditions*: Next we want to study the applicability of the proposed method to partially occluded images, where both the simulated and the real occlusions were considered.

First the case of simulated occlusion was examined. The neutral expression images [Fig. 5(a)] of the 100 selected individuals were used for learning, while the smiling, angry and screaming images with simulated partial occlusions were used for testing [Fig. 5(b)–(d)]. In each face, only those subblocks without being occluded were used for recognition. The occlusion was simulated by discarding some subblocks from the test face image, i.e., setting the gray values of all pixels within a subblock to zeros. There were four modes of occlusions simulated in all by discarding the subblocks 10% each time, in the way of row by row from the bottom to top and *vice versa* [Fig. 7(b) and (c)], and column by column from left to right and *vice versa* [Fig. 7(d) and (e)], respectively. The results at every percentage

of occlusion are denoted as  $occ_h$  in accordance with [18], with  $h = \{0, 1, 2, 3, 4, 5\}$ , for  $h = 0$ , no portions of the image were occluded; for  $h = 1$ , only 10% of the images were occluded, etc.

The results of the simulated-occlusion experiment are shown in Fig. 8, where Fig. 8(a)–(d) display the corresponding results of the occlusion modes shown in Fig. 7(b)–(e), respectively. We can find that half face occlusion does not harm the performance of the recognition system much [see Fig. 8(a), (c), and (d)] except the occlusion of upper face [see Fig. 8(b)]. This observation agrees with the previous results obtained by Brunelli and Poggio [8], i.e., the use of the mouth area may lower the performance of the system. This can be explained by the fact that the lower half contains the mouth and cheeks, which can be easily affected by most facial expression variation. However, to what extent this effect affects all recognition techniques is still a problem needing to be further studied.

Next we examined the capability of the proposed method to handle real occlusions. For that purpose, two classical wearing occlusions, the sunglasses and the scarf occlusion [Fig. 5(e), (f), (k), and (l)], were studied by conducting another experiment. As usual, the neutral expression images [Fig. 5(a)] of the 100 individuals were used for training, while the occluded images [Fig. 5(e), (f), (k), and (l)] were used for testing. Here both the training images and the testing images were not preprocessed by the histogram equalization algorithm to avoid the unwanted effect of the occluded portions of the faces, and the occluded portions of the image were not used for recognition. The results are shown in Fig. 9.

It can be seen from Fig. 9 that SOM-face consistently achieved better recognition rate than the local probabilistic approach. Specifically, when the time factor is not involved [Fig. 5(e) and (f)], the SOM-face outperformed the local probabilistic approach by about 10% or higher, concerning the *top 1 match rate*. It is interesting to note that when the duplicate images [Fig. 5(k) and (l)] were presented to the system, the results of both algorithms reveal that the occlusion of the eyes area led to better recognition results than the occlusion of the mouth area as observed in [18]. This seems to be conflicting with our previous observation [Fig. 8(b)] that upper part of the face is more important than the lower part. We think, however, that this contrary is a specific phenomenon, which is mainly caused by the fact that the scarf occluded each face irregularly and, thus, made the occluded face much more difficult to be identified than the face occluded by the sunglass (see Fig. 5).

The occlusion experiments reported previously assume that the system knows what is occluded and what is not beforehand. So how good our algorithm is when such an assumption is not valid? To answer this question, we have conducted another set of experiments as follows: the neutral expression images [Fig. 5(a)] of the 100 individuals were used for learning, while the neutral, smiling, angry and screaming images for testing [Fig. 5(a)–(d)]. To simulate the occlusion, we randomly localized a square of size  $p \times p$  ( $5 < p < 50$ ) pixels in each of the four testing image, with all the pixels inside the square set to zeros as done in [18]. Such synthetic occlusions were independently conducted 100 times, and the mean and standard deviation of the results for each of the facial expression groups were shown

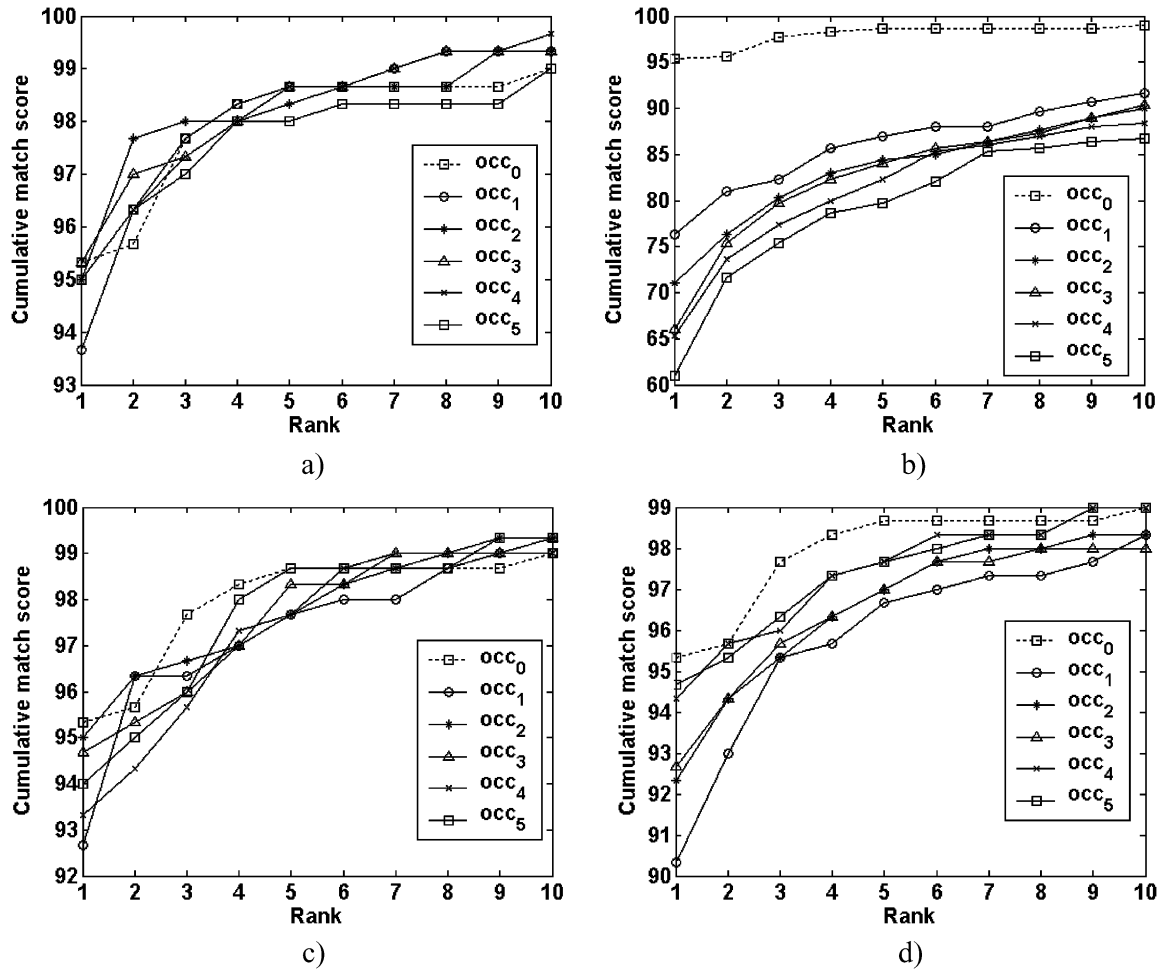


Fig. 8. Performance of the SOM-face algorithm under different occlusion modes.

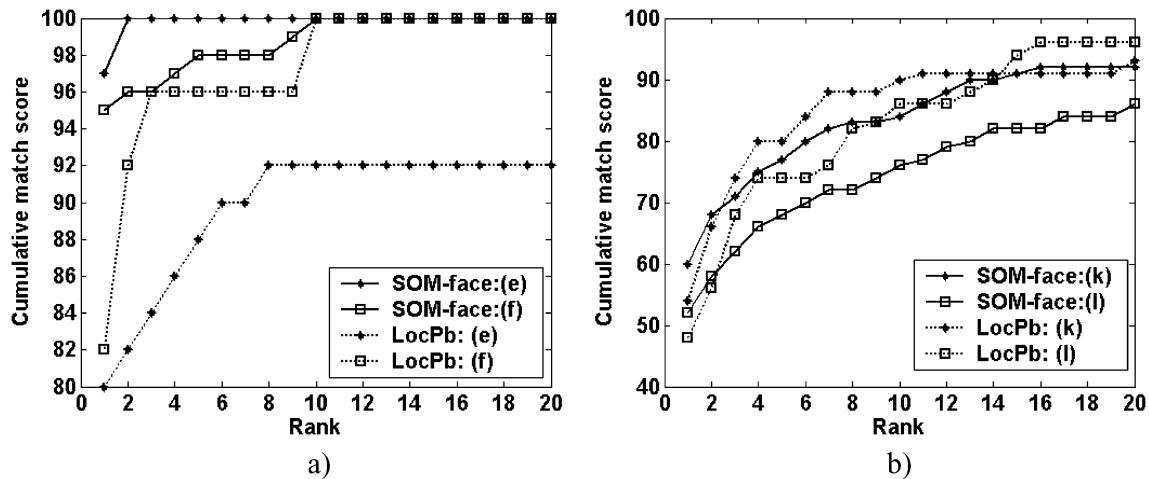


Fig. 9. Comparative performance of the SOM-face algorithm and the local probabilistic approach (LocPb) with real occlusions involved.

in Fig. 10(a), while the results one would obtain with the local probabilistic approach [18] were shown in Fig. 10(b) for comparison. Fig. 10 reveals that our method demonstrates highly robust performance against occlusions even when the occlusion size and location are unknown to the system. It can also be observed from Fig. 10 that the results on Fig. 5(a) and (b) are both so perfect that they can not be visually distinguished from the figure. Moreover, our method consistently performs better than

the local probabilistic approach, while the latter works much better than the classical PCA method [18].

3) *Visualization of the SOM-Face:* Before ending the experiments on the AR database, we will exploit the visualization capability of SOM, which enables us gain some insight into the class distribution of high-dimensional face image. In particular, the distribution of the BMUs of four images [three of which come from the same class, Fig. 11(a)–(c), one from some

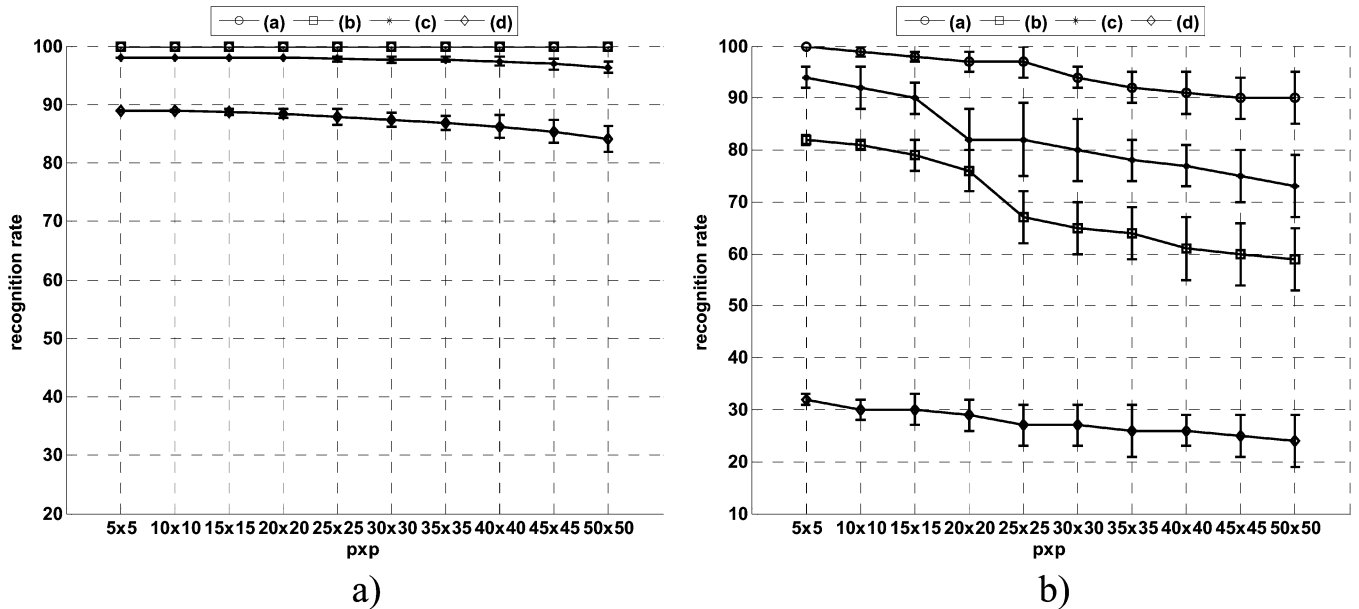


Fig. 10. Recognition rate when the occlusion size and location is unknown. (a) the SOM-face algorithm and (b) the local probabilistic approach.

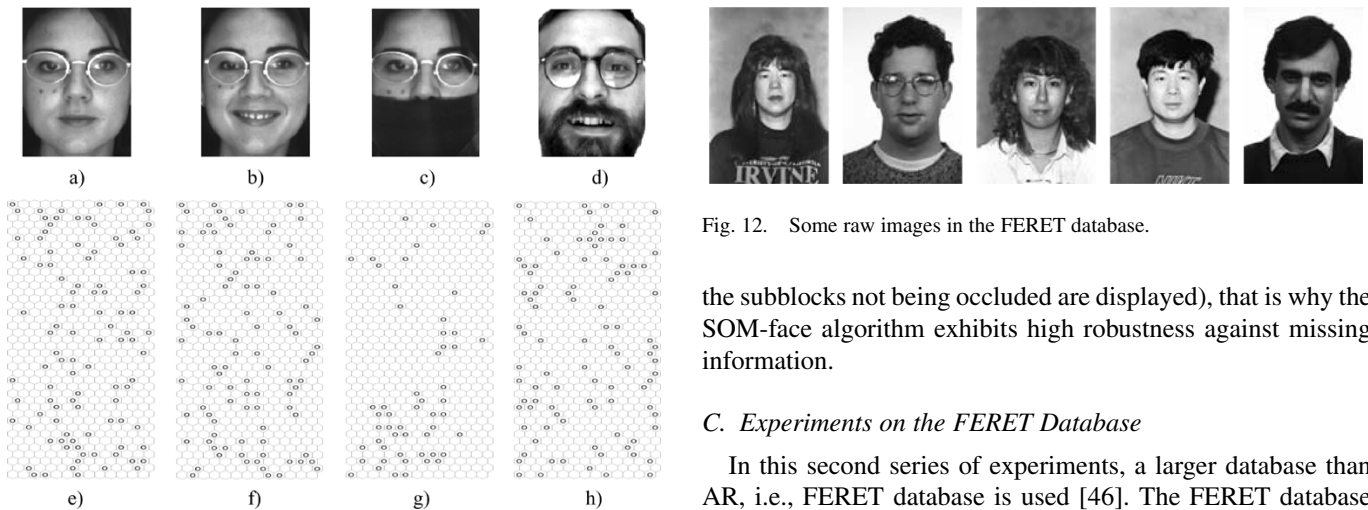


Fig. 11. Example of the corresponding distributions of the BMUs of image (a-d) in the SOM topological space (e-h). (The size of the SOM map has been scaled for better display).

other class, Fig. 11(d)] are visualized in the SOM map (Fig. 11), and the map was learned using the neutral expression images [Fig. 5(a)] of the 100 individuals, with the map size of  $12 \times 15$  neurons for better display. Several observations can be made from Fig. 11. First, the figure shows that the projections of two matching images [Fig. 11(e) and (f)] are not matched exactly in lower dimensional space. However, the similarity between the two can still be captured by the soft  $k$ -NN ensemble classifier for a correct recognition. Second, it can be observed that, in the feature space, the intra-class similarity between the two matching images [Fig. 11(e) and (f)] is larger than the inter-class similarity between the two un-matching images [Fig. 11(e) and (h)]. How to enhance such a class-specific distribution will be an interesting topic of our future research. Third, even when some of the subblocks are missing from the face (e.g., being occluded, Fig. 11(c), the left subblocks can still be used to identify the testing image (see Fig. 11(g), where only the BMUs of



Fig. 12. Some raw images in the FERET database.

the subblocks not being occluded are displayed), that is why the SOM-face algorithm exhibits high robustness against missing information.

### C. Experiments on the FERET Database

In this second series of experiments, a larger database than AR, i.e., FERET database is used [46]. The FERET database consists now of 13 539 facial images corresponding to 1565 subjects, which are diverse across gender, ethnicity, and age. Several standardized subsets of FERET images have been defined, including a common set of gallery images ( $fa$ ), and four different probe sets. In this study, the common gallery images  $fa$  are used for training, which consisted of images of 1196 people with one image per person, while  $fb$  probe set is selected for testing, which contains 1,195 images of subjects. Only facial expressions variance is involved between the corresponding images in  $fa$  and  $fb$ . Some samples of the database are shown in Fig. 12. Before the recognition process, the raw images were normalized and cropped to a size of  $60 \times 60$  pixels.

On this dataset, five experiments were conducted to further inspect some practical aspects involved in the SOM-face algorithm, such as choosing an appropriate value for the  $k$  value of the soft  $k$ -NN ensemble classifier, defining the size of subblocks, and so on. We note that in the implement of the MSOM-face strategy, alternative version of way to calculate the distance matrix is used due to its better recognition performance, i.e., combine all small maps to one big map and then calculate the distance matrix as the single SOM-face strategy.



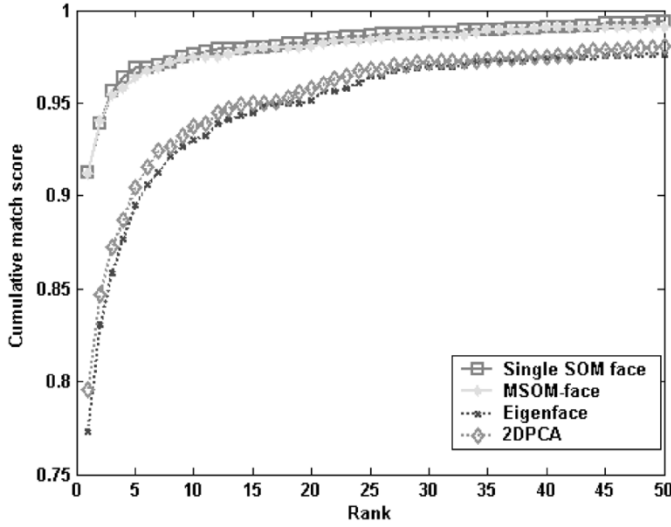


Fig. 13. Comparative performance of the single SOM-face strategy, MSOM-face strategy, eigenface and 2-DPCA algorithm on FERET.

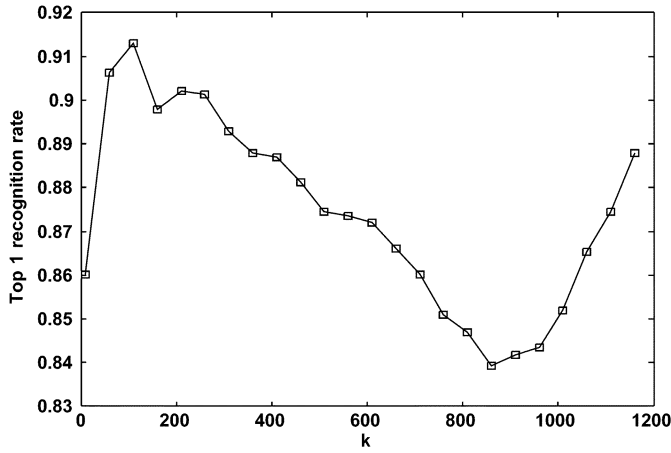


Fig. 14. Top 1 recognition rate of the single SOM-face algorithm with varying  $k$ -values.

**Experiment 1:** First, the performance of the two SOM-face based algorithms (single SOM-face and MSOM-face) on the subset was evaluated. All the images were divided into sub-blocks of  $3 \times 3$ . The results are shown in Fig. 13.

It can be observed from Fig. 13 that the results of the single SOM-face strategy and the MSOM-face strategy were close, whereas both outperformed the standard Eigenface technique and 2-DPCA by 10%–15%, respectively. Pairwise 1-tailed statistical test indicates that such difference is statistically significant at 0.95 level of significance.

**Experiment 2:** Second, we investigated the problem of choosing an appropriate  $k$ -value for the soft  $k$ -NN classifier. For this purpose, a subblock size of  $3 \times 3$  was used and then a single SOM map was trained using all 1196 images in the gallery set. And the top 1 recognition rate of the system on 1195 probe images is measured, using the soft  $k$ -NN ensemble classifier with varying  $k$ -values. The results are shown in Fig. 14. The figure reveals that it would be appropriate to set the  $k$ -value in the range of 10% to 20% of the size of the database, while the larger  $k$ -value did not necessarily produce better results. Nevertheless, as an exception to the previous remark, the largest  $k$ -value that could be set, i.e., the size of the training samples, might be an interesting point worthy to be examined.

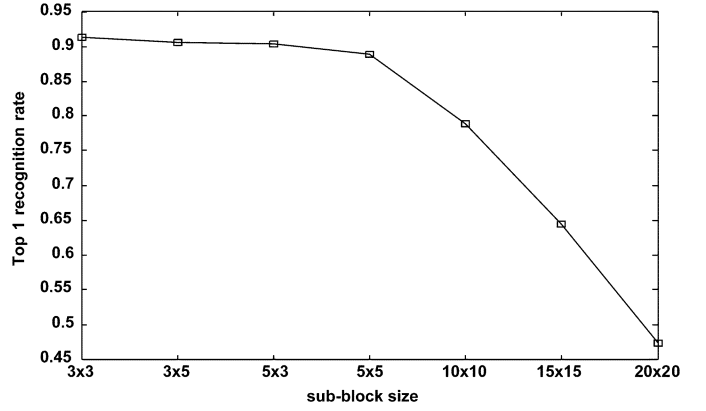


Fig. 15. Top 1 recognition rate of the single SOM-face algorithm with varying subblock size.

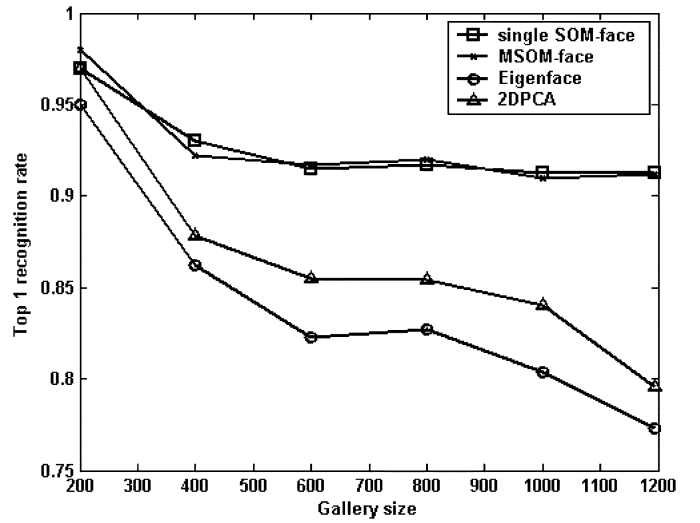


Fig. 16. Top 1 recognition rate of the single SOM-face, MSOM-face, eigenface and 2-DPCA as a function of gallery size.

**Experiment 3:** The effect of different subblock sizes on the performance of the proposed algorithm was also studied. We partitioned the training images and testing images to various subblock sizes (e.g.,  $3 \times 3$ ,  $3 \times 5$ , etc.) and, then, repeated the previous experiment (i.e., Experiment 2, but with fixed  $k$ -value) under different subblock sizes. The results are displayed in Fig. 15. It can be seen that smaller subblock seems be more beneficial to the performance of the system. Intuitively, the choice of subblock size reflects the balance between generalization and specialization, that is, as the subblocks gets smaller, the degree of generalization grows higher, while the degree of specialization becomes lower.

**Experiment 4:** The computational complexity is considered then. Formally, let the number of training faces and testing faces be  $n_{tr}$  and  $n_{test}$ , respectively. Suppose that each face is divided into  $M$  subblocks with  $d$  dimensions each subblock. Also let the number of iterations be  $T$ , the number of neurons in the SOM be  $n_c$ , and the number of neurons used in the neighborhood calculation be  $n_h$ , respectively. Then the computational complexity of the single SOM-face strategy consists of two terms, i.e.,  $O(n_{tr} M d n_h n_c T) + O(n_{test} M k d)$ , where the first term stems from the training of SOM map [39], while the second term results from the soft  $k$ -NN ensemble classifier.

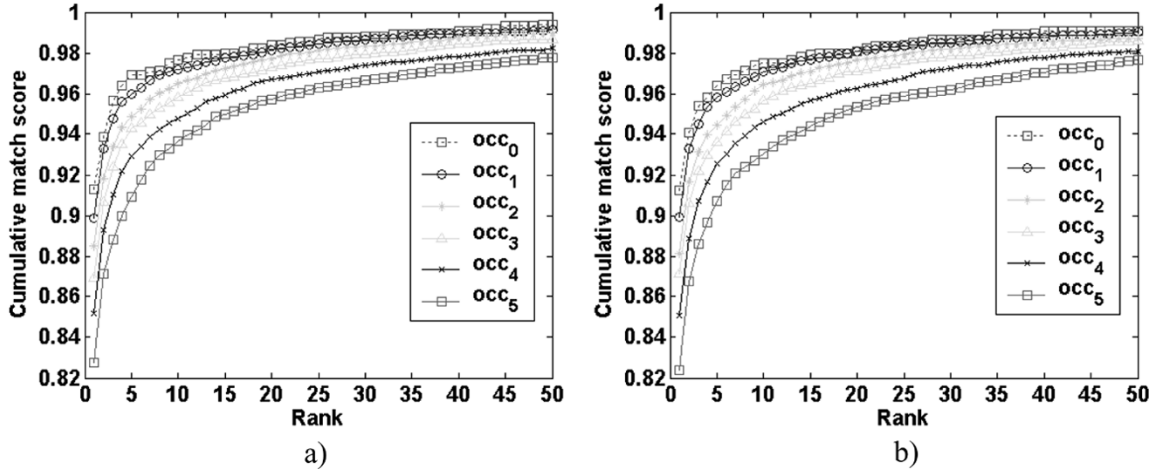


Fig. 17. Summaries of the performance of the SOM-face algorithm at different percentage of occlusions, (a) single SOM-face algorithm (b) MSOM-face algorithm.

The computational complexity of the multiple SOM-face strategy can be calculated in the same way as the single SOM-face strategy, except that a third term will be added, i.e.,  $O(n_{tr}Mn_c d)$ , which results from the recalculating of the prototype vectors when a new face is added into the database. However, the training of the MSOM-face is much faster than that of the single SOM-face due to its smaller training set and smaller SOM map. On our IBM xSeries 235 server with two Intel Xeon processors and 1GB memory, the time to train a single map of 3472 neurons for 1196 faces, with  $3 \times 3$  subblock and 500 iterations, took about 36.1 hours. However, only 1.6 hours was needed if the MSOM-face strategy was used to train 1196 small maps with 100 neurons each.

To investigate the incremental learning capability of the MSOM strategy, another experiment was conducted using different gallery sizes. Fig. 16 shows the top 1 recognition rates of different methods as the function of gallery size.

**Experiment 5:** Finally, we repeated one of the simulated occlusion experiments done on the AR dataset (see Fig. 7). In this experiment, all the images in gallery set were used for learning, while all the images in *fb* probe set were artificially partially occluded according to the same way as described in Section IV-B.2 (Fig. 7). The subblocks size was taken as  $3 \times 3$  and both single SOM-face and MSOM-face strategies were tried. Only the averaged results at different percentages of occlusions are displayed in Fig. 17. Once again, the results illustrate the robustness of the proposed method against partial occlusions.

## V. CONCLUSION

In this paper, we introduce a very simple but effective method called “SOM-face” to address the problem of face recognition with one training image per person. The images were allowed to vary in expressions and have partial occlusion. The proposed method has several advantages over some of the previous methods such as weighted local probabilistic method [18], the standard Eigenfaces technique [27] and the 2-DPCA algorithm [43]. First, it can achieve comparable or better performance than the mentioned methods with no single extra virtual samples needed to be generated. Second, it shows higher robustness

against expression variance and partial occlusions. Third, this method is very intuitive due to the visualization capability of SOM, which enables us gain some insight into the class distribution of high-dimensional face image.

We attribute these advantages to the seamless connection between the three parts of the method: By localization of the training samples, the robust performance against global changes is increased. By the unsupervised and nonparametric learning of the SOM, the similarity relationship of the subblocks in the input space is preserved in the SOM topological space. And finally, by the soft  $k$ -NN ensemble classifier, the similarity relationship is captured and exploited to enhance the whole system’s robustness.

However, it is worth mentioning that the proposed method assumes that the system knows what is occluded and what is not occluded in advance. It is natural to ask the question how to decide the location and size of the occlusion. In a general case, this is a very difficult problem since any facial portions can be occluded in any shape at any grey level. One possible way to deal with this problem is to train several specific feature detectors corresponding to each facial part (e.g., eyes, nose, mouth, and profile) [7] so as to detect the nonoccluded facial parts and use them for recognition purpose. Wu *et al.* [6] have recently presented a method to automatically locate the region of eyeglasses if the system is given a face wearing eyeglasses, using an offline trained eye region detector. Other methods such as active appearance models (AAMs) [5], deformable shape models [41] are also useful. Another way to deal with the problem is to try to bypass the problem with the help of users, for example, Gutta *et al.* discarded the whole half of the occluded face and only used the information from either left or right half of the test face for recognition [42]. Since the first way is not so mature and heavily dependent on the accuracy of the specific feature detectors, we prefer to use the second method currently, i.e., incorporating man in loop. Although experiments show that our method demonstrates highly robust performance against occlusion even when the occlusion size and location is unknown to the system, the possible need of the manual effort of the user can be regarded as a drawback of the proposed system. Thus, further studies along the first line will be the focus of our future work.

Finally, the proposed method could also be regarded as a general paradigm for dealing with *small sample problem*, in which the training set is transformed and enlarged by being partitioned into multiple subblocks. This paper shows that this paradigm works well in the scenario of face recognition with one training image per person. It is anticipated that this method is also effective in scenarios where each person has two (or more, but still *small sample*) training images.

#### ACKNOWLEDGMENT

The authors would like to thank the anonymous reviewers whose valuable comments and suggestions greatly improved this paper. They would also like to thank Dr. A. M. Martinez for providing the AR database. Portions of the research in this paper use the FERET database of facial images collected under the FERET program.

#### REFERENCES

- [1] R. Chellappa, C. L. Wilson, and S. Sirohey, "Human and machine recognition of faces: A survey," *Proc. IEEE*, vol. 83, no. 5, pp. 705–740, May 1995.
- [2] J. Daugman, "Face and gesture recognition: Overview," *IEEE Trans. Pattern Anal. Mach. Intell.*, vol. 19, no. 7, pp. 675–676, Jul. 1997.
- [3] W. Zhao, R. Chellappa, P. J. Phillips, and A. Rosenfeld, "Face recognition: A literature survey," *ACM Comput. Surv.*, pp. 399–458, Dec. 2003.
- [4] A. K. Jain, Y. Zhong, and S. Lakshmanan, "Object matching using deformable templates," *IEEE Trans. Pattern Anal. Mach. Intell.*, vol. 22, no. 1, pp. 4–37, Jan. 2000.
- [5] T. F. Cootes, G. J. Edwards, and C. J. Taylor, "Active appearance models," in *Proc. Eur. Conf. Computer Vision*, vol. 2, 1998, pp. 484–498.
- [6] C. W. C. Liu *et al.*, "Automatic eyeglasses removal from face image," *IEEE Trans. Pattern Anal. Mach. Intell.*, vol. 26, no. 3, pp. 322–336, Mar. 2004.
- [7] H. A. Rowley, S. Baluja, and T. Kanade, "Neural network-based face detection," *IEEE Trans. Pattern Anal. Mach. Intell.*, vol. 20, no. 1, Jan. 1998.
- [8] R. Brunelli and T. Poggio, "Face recognition: Features versus templates," *IEEE Trans. Pattern Anal. Mach. Intell.*, vol. 15, no. 10, pp. 1042–1062, 1993.
- [9] A. J. O'Toole and H. Abdi, "Low-dimensional representation of faces in higher dimensions of the face space," *J. Opt. Soc. Amer.*, vol. 10, no. 3, pp. 405–411, 1993.
- [10] A. K. Jain, R. P. W. Duin, and J. Mao, "Statistical pattern recognition: A review," *IEEE Trans. Pattern Anal. Mach. Intell.*, vol. 22, no. 1, pp. 4–37, Jan. 2000.
- [11] S. C. Chen, J. Liu, and Z.-H. Zhou, "Making FLDA applicable to face recognition with one sample per person," *Pattern Recognit.*, vol. 37, no. 7, pp. 1553–1555, 2004.
- [12] J. Lu, K. N. Plataniotis, and A. N. Venetsanopoulos, "Face recognition using Kernel direct discriminant analysis algorithms," *IEEE Trans. Neural Netw.*, vol. 14, no. 1, pp. 117–126, 2003.
- [13] D. L. Swets and J. Weng, "Using discriminant eigenfeatures for image retrieval," *IEEE Trans. Pattern Anal. Mach. Intell.*, vol. 18, no. 8, pp. 831–836, Aug. 1996.
- [14] P. N. Belhumeur, J. P. Hespanha, and D. J. Kriegman, "Eigenfaces versus Fisherfaces: Recognition using class specific linear projection," *IEEE Trans. Pattern Anal. Mach. Intell.*, vol. 19, no. 7, pp. 711–720, Jul. 1997.
- [15] I. Lapidot, H. Guterman, and A. Cohen, "Unsupervised speaker recognition based on competition between self-organizing maps," *IEEE Trans. Neural Netw.*, vol. 13, no. 4, pp. 877–887, Jul. 2002.
- [16] B. Moghaddam, T. Jebara, and A. Pentland, "Probabilistic visual learning for object detection," in *Proc. Int. Conf. Computer Vision*, 1995, pp. 786–793.
- [17] A. M. Martinez and A. C. Kak, "PCA versus LDA," *IEEE Trans. Pattern Anal. Mach. Intell.*, vol. 23, no. 2, pp. 228–233, Feb. 2001.
- [18] A. M. Martinez, "Recognizing imprecisely localized, partially occluded, and expression variant faces from a single sample per class," *IEEE Trans. Pattern Anal. Mach. Intell.*, vol. 25, no. 6, pp. 748–763, Jun. 2002.
- [19] S. Chen, D. Zhang, and Z.-H. Zhou, "Enhanced (PC)<sup>2</sup>A for face recognition with one training image per person," *Pattern Recognit. Lett.*, vol. 25, pp. 1173–1181, 2004.
- [20] D. Beymer and T. Poggio, "Face recognition from one example view," *Science*, vol. 272, no. 5250, 1996.
- [21] A. M. Martinez, "Recognition of partially occluded and/or imprecisely localized faces using a probabilistic approach," *Proc. IEEE Computer Vision and Pattern Recognition (CVPR)*, vol. 1, pp. 712–717, 2000.
- [22] A. M. Martinez, "Matching expression variant faces," *Vis. Res.*, vol. 43, no. 9, pp. 1047–1060, 2003.
- [23] S.-H. Lin, S. Y. Kung, and L.-J. Lin, "Face recognition/detection by probabilistic decision-based neural network," *IEEE Trans. Neural Netw.*, vol. 8, no. 1, pp. 114–132, Jan. 1997.
- [24] T. Kohonen, *Self-Organization and Associative Memory*, 2nd ed. Berlin, Germany: Springer-Verlag, 1988.
- [25] —, *Self-Organizing Map*, 2nd ed. Berlin, Germany: Springer-Verlag, 1997.
- [26] B. Zhang, H. Zhang, and S. Ge, "Face recognition by applying wavelet subband representation and Kernel associative memory," *IEEE Trans. Neural Netw.*, vol. 15, no. 1, pp. 166–177, 2004.
- [27] L. Sirovich and M. Kirby, "Low dimensional procedure for the characterization of human faces," *JOSA-A(4)*, no. 3, pp. 519–524, 1987.
- [28] A. Pentland, B. Moghaddam, and T. Starner, "View-based and modular eigenspaces for face recognition," in *Proc. IEEE Int. Conf. Computer Vision and Pattern Recognition*, Seattle, WA, 1994, pp. 84–91.
- [29] D. Valentin *et al.*, "Connectionist models of face processing: A survey," *Pattern Recognit.*, vol. 27, no. 9, pp. 1209–1230, 1994.
- [30] P. R. Villela and J. van Humberto Sossa Azucla, *Improving Pattern Recognition Using Several Feature Vectors*. Berlin, Germany: Springer-Verlag, 2002, Lecture Notes in Computer Science.
- [31] L. I. Kuncheva and C. J. Whitaker, *Feature Subsets for Classifier Combination: An Enumerative Experiment*, J. Kittler and F. Roli, Eds. Berlin, Germany: Springer-Verlag, 2001, vol. 2096, Lecture Notes in Computer Science, pp. 228–237.
- [32] Y. Linde, A. Buzo, and R. M. Gray, "An algorithm for vector quantizer design," *IEEE Trans. Commun.*, vol. COM-28, pp. 84–95, Jan. 1980.
- [33] W. Pedrycz, "Fuzzy sets in pattern recognition: Methodology and methods," *Pattern Recognit.*, vol. 23, no. 1/2, pp. 121–146, 1990.
- [34] S. Lawrence, C. L. Giles, A. Tsoi, and A. Back, "Face recognition: A convolutional neural-network approach," *IEEE Trans. Neural Netw.*, vol. 8, no. 1, pp. 98–113, Jan. 1997.
- [35] C. Liu and H. Wechsler, "Evolutionary pursuit and its application to face recognition," *IEEE Trans. Pattern Anal. Mach. Intell.*, vol. 22, no. 6, pp. 570–582, Jun. 2000.
- [36] L. Wiskott, J. M. Fellous, N. Kruger, and C. V. D. Malsburg *et al.*, "Face recognition by elastic bunch graph matching," in *Intelligent Biometric Techniques in Fingerprint and Face Recognition*, L. C. Jain *et al.*, Eds. Berlin, Germany: Springer-Verlag, 1999.
- [37] M. Fleming and G. Cottrell, "Categorization of faces using unsupervised feature extraction," in *Proc. IEEE Int. Joint Conf. Neural Networks*, vol. 2, 1990, pp. 65–70.
- [38] J. D. McAuliffe, L. E. Atlas, and C. Rivera, "A comparison of the LBG algorithm and Kohonen neural network paradigm for image vector quantization," in *Proc. Int. Conf. Acoustics, Speech and Signal Processing*, vol. IV, 1990, pp. 2293–2296.
- [39] T. Kohonen *et al.*, "Self organization of a massive document collection," *IEEE Trans. Neural Netw.*, vol. 11, no. 3, May 2000.
- [40] J. A. Kangas, T. Kohonen, and J. T. Laaksonen, "Variants of self-organizing maps," *IEEE Trans. Neural Netw.*, vol. 1, no. 1, pp. 93–99, Jan. 1990.
- [41] T. K. Leung, M. C. Burl, and P. Perona, "Finding faces in cluttered scenes using random labeled graph matching," in *Proc. 5th IEEE Int. Conf. Computer Vision*, 1995, pp. 637–644.
- [42] S. Gutta and H. Wechsler, "Face recognition using asymmetric faces," in *Proc. 1st Int. Conf. Biometric Authentication (ICBA)*, D. Zhang and A. K. Jain, Eds., 2004, vol. 3072, pp. 162–168.
- [43] J. Yang, D. Zhang, A. F. Frangi, and J. Yang, "Two-dimensional PCA: A new approach to appearance-based face representation and recognition," *IEEE Trans. Pattern Anal. Mach. Intell.*, vol. 26, no. 1, pp. 131–137, Jan. 2004.
- [44] A. M. Martinez and R. Benavente, "The AR Face Database," CVC Technical Report, no. 24, 1998.
- [45] —, (2003) The AR Face Database. [Online]. Available: [http://rv11.ecn.purdue.edu/~aleix/aleix\\_face\\_DB.html](http://rv11.ecn.purdue.edu/~aleix/aleix_face_DB.html)
- [46] P. J. Phillips, H. Wechsler, J. Huang, and P. J. Rauss, "The FERET database and evaluation procedure for facerecognition algorithms," *Image Vis. Comput.*, vol. 16, no. 5, pp. 295–306, 1998.

- [47] W. Yambor, B. Draper, and R. Beveridge, "Analyzing PCA-based face recognition algorithms: Ei-genvektor selection and distance measures," in *Empirical Evaluation Methods in Computer Vision*, H. Christensen and J. Phillips, Eds. Singapore: World Scientific, 2002.
- [48] P. S. Penev and J. J. Atick, "Local feature analysis: A general statistical theory for object representation," *Network: Computat. Neural Syst.*, vol. 7, no. 3, pp. 477–500, 1996.
- [49] J. Kittler, M. Hatef, R. Duin, and J. Matas, "On combining classifiers," *IEEE Trans. Pattern Anal. Mach. Intell.*, vol. 20, no. 3, pp. 226–239, Mar. 1998.
- [50] R. Fergus, P. Perona, and A. Zisserman, "Object class recognition by unsupervised scale-invariant learning," in *Proc. Int. Conf. Computer Vision Pattern Recognition*, vol. 2, Madison, WI, Jun. 2003, pp. 264–275.



**Xiaoyang Tan** received the B.Sc. and M.Sc. degrees in computer science from Nanjing University of Aeronautics and Astronautics, Nanjing, China, in 1993 and 1996, respectively. He is currently working toward the Ph.D. degree in the Department of Computer Science and Technology, Nanjing University.

His research interests are in machine learning, pattern recognition, and computer vision.



**Songcan Chen** received the B.Sc. degree in mathematics from Hangzhou University (now merged into Zhejiang University) in 1983, the M.Sc. degree in computer applications from Shanghai Jiaotong University, Shanghai, and the Ph.D. degree in communication and information systems from Nanjing University of Aeronautics and Astronautics (NUAA), Nanjing, China, in 1997.

He has been a Full Professor in the Department of Computer Science and Engineering, NUAA, since 1998. His research interests include pattern recognition, machine learning, and neural computing. In these fields, he has authored or coauthored over 70 scientific journal papers.

tion, machine learning, and neural computing. In these fields, he has authored or coauthored over 70 scientific journal papers.



**Zhi-Hua Zhou** (S'00–M'01) received the B.Sc., M.Sc., and Ph.D. degrees in computer science from Nanjing University, Nanjing, China, in 1996, 1998, and 2000, respectively, all with the highest honor.

He joined the Department of Computer Science and Technology, Nanjing University, as a Lecturer in 2001, and is currently a Professor and Head of the LAMDA group at present. His research interests are in artificial intelligence, machine learning, data mining, pattern recognition, information retrieval, neural computing, and evolutionary computing. In

these areas he has published over 40 technical papers in refereed international journals or conference proceedings. He is an Associate Editor of *Knowledge and Information Systems* and is on the Editorial Boards of *Artificial Intelligence in Medicine* and the *International Journal of Data Warehousing and Mining*.

Dr. Zhou won the Microsoft Fellowship Award in 1999, the National Excellent Doctoral Dissertation Award of China in 2003, and the Award of National Outstanding Youth Foundation of China in 2004. He served as the organizing Chair of the 7th Chinese Workshop on Machine Learning in 2000, was program Co-chair of the 9th Chinese Conference on Machine Learning in 2004, and Program Committee Member for numerous international conferences. He is a senior member of China Computer Federation (CCF) and the Vice-chair of the CCF Artificial Intelligence and Pattern Recognition Society, a Councilor of Chinese Association of Artificial Intelligence (CAAI), the Vice-chair and Chief Secretary of CAAI Machine Learning Society, and a Member the IEEE Computer Society.



**Fuyan Zhang** is a Professor in the Department of Computer Science and Technology, Nanjing University, Nanjing, China. His research interests are in information processing and multimedia. In these areas, he has published over 100 technical papers in refereed international journals or conference proceedings. He is currently the Chair of the Multimedia Computing Institute of Nanjing University.