

# Human recognition based on ear shape images using PCA-Wavelets and different classification methods

Ali Mahmoud Mayya1\* and Mariam Mohammad Saii

PhD student, Computer Engineering, Tishreen University, Syria

## Abstract

A new approach of human recognition using ear images is introduced. It consists of two basic steps which are the ear segmentation and ear recognition. In the first one, Likelihood skin detector is used to determine the skin areas in the side face images. Then, some of the morphological operations are applied to determine the ear region. This region is extracted using image processing techniques. The ear recognition step depends on the segmented ear images as inputs. A hybrid PCA\_Wavelet algorithm is used to extract the ear features from ear. Finally, the feed forwarding back propagation neural network is trained using the feature vectors. Tests which applied on 460 images, which have been taken during 4 months and under different illumination and pose variations, show that the system achieved a rate of 96.73% for ear extraction and 98.9% for recognition. More experiments are done to specify the best wavelet level, the best number of features, the best classification method, and the best threshold value. The study is also compared with other ones at the area of ear recognition.

## Introduction

Ear recognition is an example of the Human recognition using biometrics depending on the human biological properties. This kind of recognition has been recently used because of the ear's distinctive properties such as its invariant shape.

The ear is defined by unique shape that contains different parts (in location and size), and obvious curved lines that facilitate the feature extraction phase. Even the ears of "identical twins" differ in some respects [3,7,15,16], also Ear is less susceptible to distortions than the fingerprint or handprint (fingerprint may suffer from burns, wounds or other deformities) so, ear is used for recognition especially in air planes accidents [5,14]. In addition to that, ear images is easier to be collected than any other biometrics since it doesn't require any tools to be inserted in human body to collect samples as in DNA, and this is why ear is considered non-invasive. We can also observe that ear images are smaller than face, hand or leg ones which makes the recognition process fast and easy [14]. The most important thing is that ear maintains its shape between 8 and 70 of ages so, its changes are very slow and little [11].

## Related work

A. Ianarelli [10] was the first actual person who used ear to recognize people. He compared 10,000 ears drawn from a randomly selected sample in California.

Although his study was important, it stills difficult to apply on computer because of the difficulty of localizing the anatomical point which serves as the origin of the measurement of the system [1].

Burge and Burger [1] proposed the use of Voronoi diagrams. Their technique used an adjacency graph built from Voroni regions of ear-curve segments. Although their paper had no recognition results, it prompted a range of further studies into the effectiveness of ears as a biometric [2]. Their system didn't take in account illumination and pose variation.

Morone *et al.* [13] tried three neural networks approaches to recognize people based on ear images. These approaches were (Borda, Bayesian, and Weighted Bayesian combination) which depended on the macrofeatures extracted from the ear. As a result, they got 93% recognition rate.

The genetic search was a powerful approach used by Yuizono [18] to design a robust ear recognition system. Automatic extraction using template matching was used to obtain ear images. The proposed system included 6 images from video frame for each individual; the first three images used for database while the second three images used as test images. No illumination or pose variations were taken in account. Occlusion by hair or earrings also was not considered. The recognition rate was 100%.

Later, Hurley [9] applied the force field transform and PCA to extract the force features like energy lines, in the classification stage they used the Euclidian distance classifier to obtain 99.2% recognition rate.

Mahbubur Rahman [15] used the Generalized Hugh Transform (GHT) to extract the ear's features from the ear image which was extracted from the acquired image using Masking. The extracted features along with the subjects' id were stored in the database for testing for a match. As a result they got almost 89% recognition rate.

Recently, Islam [11] had proposed using AdaBoost algorithm to extract ear images. They also trained their system with rectangular haar-like features depending on a set of training images with different

**Correspondence to:** Ali Mahmoud Mayya, Computer and Automatic Control Engineering department, Tishreen University, Syria, E-mail: alimia1988@yahoo.com

**Key words:** ear recognition, ear segmentation, ear images, back propagation, PCA, Wavelet

**Received:** August 14, 2016; **Accepted:** September 12, 2016; **Published:** September 16, 2016

age and gender and under various illumination and pose conditions. Later, they updated their system by means of ICP algorithm. They obtained a rank one recognition rate of 93% while testing with the University of Notre Dame Biometrics Database.

Zhichun Mu and Zhaoxia Xie [17] compared the use of locally linear embedding (LLE) and PCA algorithms; they found that the first algorithm is better in recognition and when they applied the improved version IDLLE, the recognition rate increased precisely. They also found that recognition decreased with pose variation, but they commonly got 80% recognition rate.

In a recent publication of Kumar and Wu [12] they present an ear recognition approach, which uses the phase information of Log-Gabor filters for encoding the local structure of the ear. The encoded phase information is stored in normalized grey level images. The rank-e performance for the Log-Gabor approaches ranges between 92.06% and 95.93% on a database which contains 753 images from 221 subjects.

## Our work

The proposed ear recognition system consists of two main parts shown in Figure 1. The first part, ear segmentation, includes three steps which are the skin detection, morphological operations and ear extraction. The second part, ear recognition, includes the feature extraction steps (Wavelet approximation plus PCA) and training the neural network classifier. Each part of these steps is described in the following sections.

## Ear segmentation part

The input of this step is the side face images which will be segmented using the likelihood skin detector and the morphological operations.

### Skin detection

In this stage, each pixel of side face image is studied individually to be classified to either skin or non-skin pixels depending on the likelihood ratio. We use a skin detector developed by Ciarán Ó Conaire. [6] that calculates the most likely skin pixels based on a previously computed skin model. Their non-parametric histogram-based models were trained using manually-annotated skin pixels (14,985,845 pixels) and non-skin pixels (304,844,751 pixels).

Image pixels are classified according to the value  $P(\text{Skin}|C)$  which represents the probability that the pixel whose color is  $C$  belongs to the skin region. Because the computations of all probabilities are not possible, the Bayes rule of computing probability of observing skin, given  $c$  color  $P(\text{Skin}|C)$  is given as follows [8]:

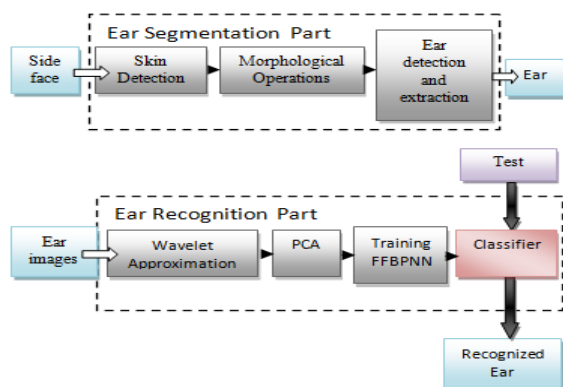


Figure 1. Block diagram of the proposed ear recognition system.

$$P(\text{Skin}|C) = \frac{P(C|\text{Skin})P(\text{Skin})}{P(C|\text{Skin})P(\text{Skin}) + P(C|N-\text{Skin})P(N-\text{Skin})}$$

After computing the posterior probability  $P(\text{Skin}|C)$ , it must be compared with  $\theta$  which is defined as follows:

$$0 \leq \theta \leq 1$$

If this probability is greater than  $\theta$ , the pixel is considered as a skin pixel; on other hand, the pixel is labeled as non-skin pixel. After obtaining the likelihood images, these images are transformed to the binary format.

### Ear detection

This stage uses the morphological operations to process the resultant images from the previous stage which contain many undesired white and black points which must be deleted. To achieve this, two morphological operations are done. The first is an operation of filling holes which is performed to remove the black points inside the white regions. This operation produces an assembled white region representing the ear area.

However; if the resultant image still contains some noisy white points on the borders, it will be removed by clearing borders operation. The second operation is the erosion operation which is done to decrease the white areas in the image to obtain the right ear region. Figure 2 illustrates the detailed skin and ear detection phases explained previously.

### Ear extraction

The proposed approach of obtaining the ear region depends on cropping this region from the original side face image. So, it needs the start point coordinates ( $x_{min}$ ,  $y_{min}$ ), the width, and the height of the cropped rectangle which will be obtained from the original image by one of two proposed ways. The first way is experimental one, which requires vertical scanning for the first white pixel in the image. Then, the width and height are defined experimentally because the distance between side face and camera is almost 20-25 cm and the ear region mostly resides at the left angle of the image. The second way, the measurement-based way, finds four basic points illustrated in Figure 3A. First, an up-down scanning is done to identify the start point (first white pixel), and the end point (last white pixel). Second, a scanning

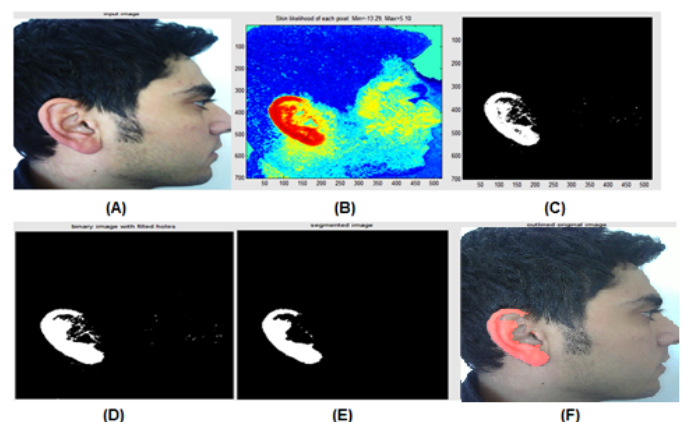


Figure 2. Ear detection phases.

A: original image, B: likelihood image of image (A), C: binary image of image (B), D: filled holes image of image (C), E: (D) image after erosion, F: output image (image with colored ear region)

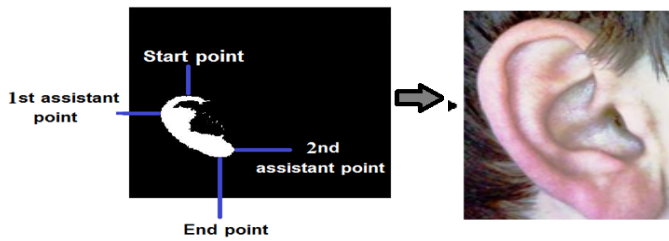


Figure 3. Ear extraction: (A) The measurement method (B) The output image.

from left to right is performed to identify the first assistant point (first white pixel) and the second assistant point (last white point). The width and height can be defined by means of simple subtraction. For more understanding, the following example is mentioned:

Suppose that 1st point is:  $s(m1,n1)$ ; 2nd point is:  $a1(m2,n2)$ ; 3rd point  $a2(m3,n3)$ ; 4th point  $e(m4,n4)$ , then the cropping rectangle will have the upper-left point  $(n2,m1)$ ,  $n3-n2$  width and  $m4-m1$  height.

Figure 3 illustrates the way of defining the rectangular area which will contain the ear region based on the measurement-based method.

### Ear recognition part

This step depends on the output ear images of the previous part. These images are used as input to the feature extraction step.

### Feature extraction

Feature extraction of the ear images is very significant stage to obtain the most important information of original images in order to reduce the next step's processing time.

There are different algorithms which can be used to extract features, but the proposed approach suggests using wavelet transform to obtain the approximation coefficients (cA) which will be normalized to match the [0-1] form. This normalized (cA) image will be stored as (.jpg) image, and this will make its gray values in the range [0-255]. Next, principle component analysis algorithm is used to get the final feature vector of each image.

**Wavelet 2D transform:** The wavelet transform concentrates the energy of the image into a small number of wavelet coefficients.

The importance of this step is to minimize the size of the original ear images, and extract the most important values of them. In this study, the 2D wavelet transform of level 2 is used to decompose the ear image into coefficients. Then, the approximation coefficient (cA) is normalized by dividing its values by maximum of (cA). This process insures that the maximum value of this matrix will be 1, and the pixel's range is [0-1]. In the next step, the normalized image is stored as JPG format, and the gray scale range becomes [0-255]. At the end of this step, the image's size will be almost four times smaller than the original size.

Using two dimensional wavelet transforms, an image  $f(x,y)$  can be represented as follows [4]:

$$f(x,y) = S_j \phi_j(x,y) + \sum_{j=1}^J d_j^1 \phi_j^1(x,y) + \sum_{j=1}^J d_j^2 \phi_j^2(x,y) + \sum_{j=1}^J d_j^3 \phi_j^3(x,y) + S_j + \sum_{j=1}^J D_j^1 + \sum_{j=1}^J D_j^2 + \sum_{j=1}^J D_j^3 \quad (3)$$

Where the two dimensional wavelets are the tensor product of the one dimensional wavelets as below [4]:

$$\phi_j(x,y) = \phi(x) \times \phi(y) \quad (4)$$

$$\varphi^v(x,y) = \phi(x) \times \varphi(y) \quad (5)$$

$$\phi^h(x,y) = \phi(x) \times \phi(y) \quad (6)$$

$$\varphi^d(x,y) = \phi(x) \times \varphi(y) \quad (7)$$

Where J represents the number of wavelet levels. The first stage is called the approximation coefficients, where the image's energy is concentrated. The other components are called the detailed coefficients which are the horizontal, vertical and diagonal images ( $cH, cV, cD$ ).

**Principle component analysis:** After the ear images have been transformed and normalized in the previous stage, a principle component analysis (PCA) is performed to extract the final feature vectors.

PCA is a technique for reducing the dimension of feature vectors while preserving the variation in the dataset. A low dimension space called 'Eigen space', which is defined by a set of 'Eigen vectors' of dataset is used in classification. In ear biometric, the eigenvalues and eigenvectors are computed for the set of training images, and an "ear space" is selected based on the eigenvectors.

After computing the eigenvectors and eigenvalues of each approximated image, the eigenvectors corresponds to the eigenvalues that exceeds the threshold ( $0.5e+008$ ) are selected, while the others are eliminated.

This process will produce the last feature vectors.

Figure 4A shows feature vectors of two different ear images related to individuals who are twins. The plot shows the intersection between the two vectors. Although twins have similar ears, the vector features of them are little different. So, it can be said that the proposed feature extraction method are effective even in the case of twins' images.

On the other hand, Figure 4B illustrates the variance between feature vectors of two different ear images related to non-relative individuals.

### Back propagation neural networks

The proposed method suggests using feed forwarding back propagation neural networks FFBPNN as a classifier. The FFBPNN proposed topology is illustrated in Figure 5.

The chosen FFBPNN has one hidden layer consisted of 1000 neurons, 200 neurons output layer, ('tansig','logsig') functions for the hidden and output layer, and 'trainscg' as learning function as shown in Figure 5.

### Experiments and results

Experiments have been performed to evaluate the efficiency and robustness of the system using our ear database. The database contains

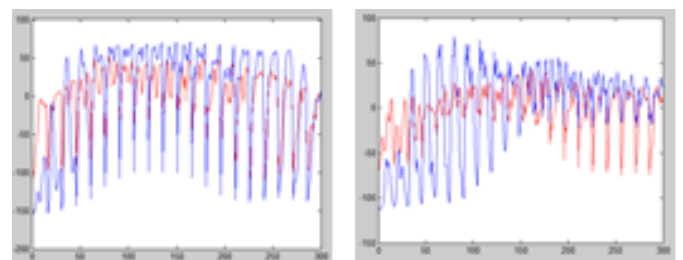


Figure 4. Feature vectors of two images: (A): relative individuals (B): non-relative individuals.

460 side face images corresponding to 55 persons of age between 12 and 60. All the illumination variations, pose variations, day variations and even night-day variations are taken in account during imaging.

### Ear segmentation tests

The system succeeded in segmenting ear image correctly from 445 side face images, but it failed in 15 images, and this achieved 96.73% rate. Figure 6 illustrates examples of the side face images succeeded in segmentation phase.

It can be noticed from Figure 6 that images have different illumination and pose variations. Some images suffer from degradation

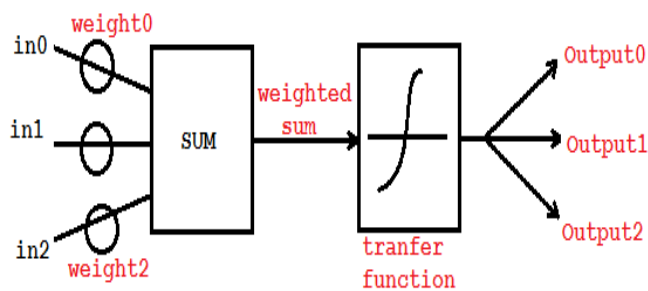


Figure 5. Neuron Model.

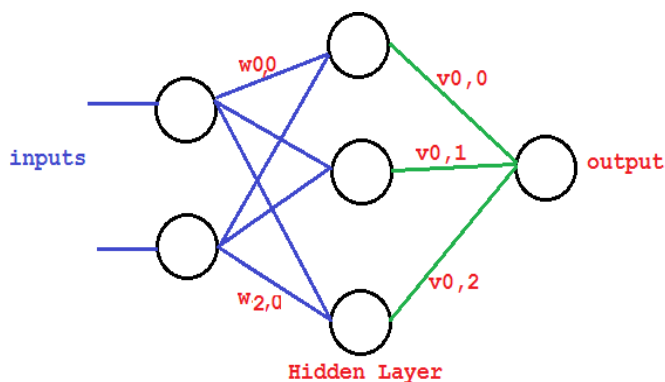


Figure 6. Neural Networks Model.

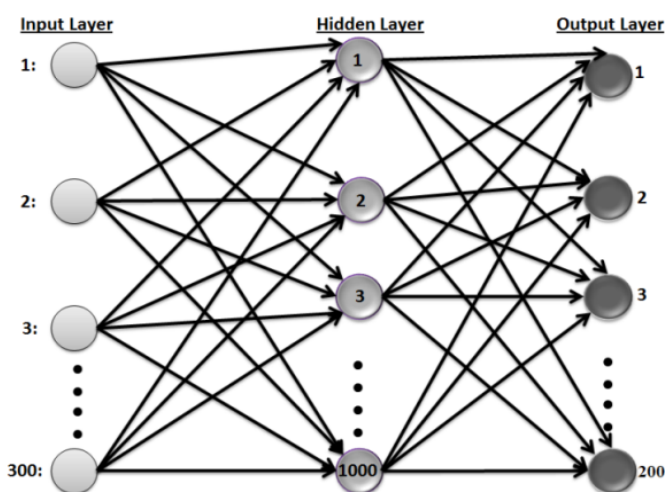


Figure 7. The proposed back propagation network topology.

due to the low level of brightness or occlusion by hair and earrings. Variable distances from camera and various side faces' skin color are taken in account. In Figure 7, there are four images which have bad ear area because of occlusion of the most ear area (A), occlusion and bad imaging (B), bad pose from camera (C,D).

### Ear recognition tests

200 samples (4 to each individual) of correctly segmented ear images were chosen to build the ear database and train FFBP classifier. The system achieved 98.9% recognition rate when it was tested on 170 ear images which don't belong to the ear database. Here are some of the ear images which recognized correctly.

Figure 8 includes different ear images which have different illumination and poses. Some of them are covered by hair or earrings, the others suffer from degradation due to camera motion or to bad imaging. All of these images are classified correctly. On the other hand, Figure 9 illustrates the four samples which recognized incorrectly, and this is because of big degradation that caused by occlusion of the most ear area (A, B), and bad segmented ear (C, D).

**Determining the best network parameters:** Experiments were performed to determine the best FFBP network to use. Figure 10 introduces two different network performance analyses describing the importance of choosing the appropriate network parameters (training function, layer function, number of layers, and number of neurons etc.).

(A) 'tansig', 'purelin', 'traingd',

(B) use 'tansig', 'logsig', 'trainscg'

In Figure 10A, it can be noticed that the performance can't match the goal (0.02) due to the selective layer and training (learning) functions which are 'tansig', 'purelin', 'traingd' respectively. In contrast, the performance is very good and matches the goal at the 24th epoch when we use 'tansig', 'logsig', 'trainscg' functions in Figure 10B, while it

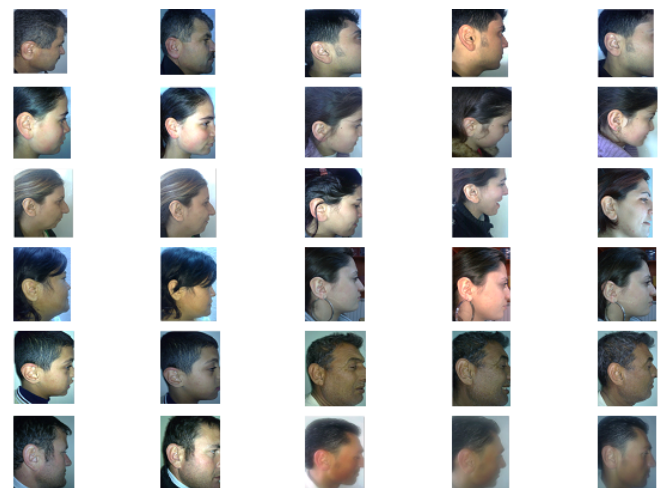


Figure 8. Examples of correctly segmented samples.



Figure 9. Example of failed samples.

spends 1000 epochs in Figure 10A with lower performance.

Another parameter which plays such important role in network performance is the number of neurons in each layer. Figure 11 explains the network's performance according to number of neurons. Figure 11A shows network's performance according to [2000 200] neurons of the hidden and output layer respectively. It can be noticed that this performance is very similar to another one in Figure 11B where network has [1000 200] form. The last option of 500 neuron of the hidden layer requires the least number of epochs, and the network's error drops fast, Figure 11C.

(A) 2000 for 1st layer, 200 for 2nd one

(B) 1000 for 1st layer, 200 for 2nd one

(C) 500 for 1st layer, 10 for 2nd one

Table 1 includes the recognition rates of different FFBP neural networks. The networks differ in number of epochs, layers, neurons and training time.

According to previous analysis, the FFBPNN with one hidden layer consisted of 1000 neurons, 200 neurons output layer, ('tansig','logsig') functions for the hidden and output layer, and 'trainscg' as learning function is the best back propagation topology to use

Figure 12 includes the performance of the two selected FFBP networks according to different numbers of epochs. It can be concluded

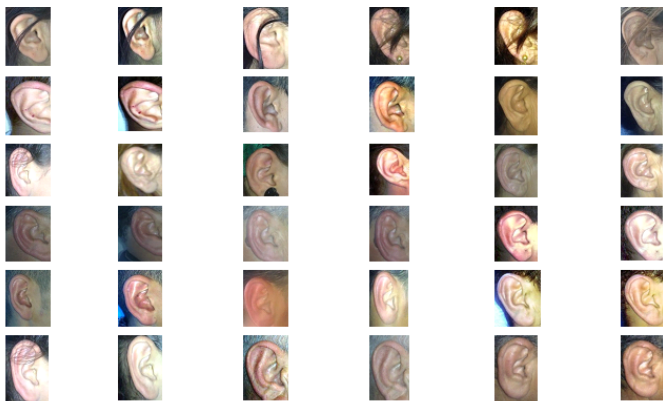


Figure 10. Examples of correctly recognized ear images.

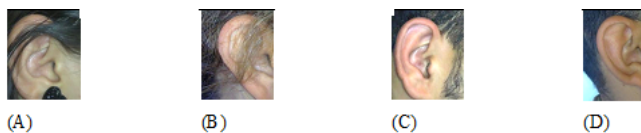


Figure 11. The incorrectly recognized ear images.

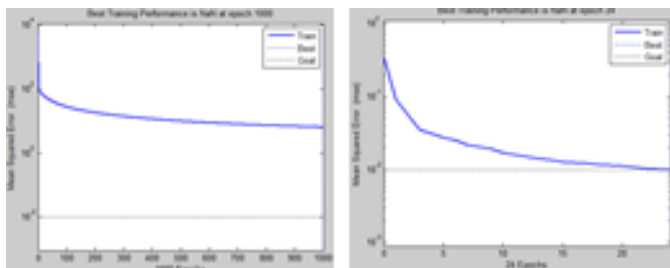


Figure 12. Performance analysis of ffbpnn according to layers and training functions.

that the second choice, the blue curve, ('tansig','logsig') is better than the first, the red one, ('tansig','purelin').

So, sections 4.2.2, 4.2.3 and 4.2.4 will be introduces in the term of the this selective topology.

**Determining the wavelet's level:** To determine the best wavelet's level, some of the experiments were done, and the back propagation classifier was evaluated by different feature vectors which have been extracted under different wavelet's levels.

From Table 2, it can be concluded that approximation components (cA) of the second level give the best recognition rate. So, the level 2 was selected.

**Choosing the appropriate threshold:** Choosing the best threshold is very important problem to make system capable of refusing the strange individuals and accept the system's persons. The number of samples that failed due to the threshold constraint must be defined, and this will help to get the best threshold. To make a best selection, 20 different samples of stranger individuals which don't belong to the system are introduced to test the system's rejection performance. Table 3 describes the number of failed samples according to each threshold value. Table 3 shows that the best threshold value with least number of false positives is 0.008, but at the same time, it can be noticed that the performance decreased to 94.32%. Figure 13 includes the false acceptance and rejection.

**Choosing the appropriate number of elements of feature vector:** We suggest using the first 30th elements of feature vector, the first 90th and the first 150th ones to test system instead of using the hall feature vectors. Table 4 includes the results of using those choices. It can be concluded that using less features may reduce the training time but led to lower performance. For more accurate analysis, we compute the mean squared error (MSE) of each test sample introduced to the

Table 1. Experimental results for different FFBPNN.

| Epochs | Layers | Neurons         | Training Time | Rate % |
|--------|--------|-----------------|---------------|--------|
| 1000   | 2      | [2000 200]      | 5:14          | 95.8   |
| 524    | 2      | [1000 200]      | 3:13          | 98.9   |
| 255    | 2      | [500 10]        | 0:36          | 98.6   |
| 152    | 2      | [1000 10]       | 0:41          | 92.4   |
| 159    | 3      | [1000 1000 200] | 2:48          | 97.7   |

Table 2. The recognition rate according to different wavelet's levels.

| Wavelet level | Feature extraction time (minutes) | FFN epochs | Rate % |
|---------------|-----------------------------------|------------|--------|
| 2             | 0.1413                            | 304        | 98.9   |
| 3             | 0.1439                            | 427        | 95.4   |
| 4             | 0.1589                            | 256        | 92.1   |

Table 3: Experimental results of changing threshold values.

| Threshold | False Rejection | False Acceptance | Rate % |
|-----------|-----------------|------------------|--------|
| 0.122     | 1               | 20               | 98.6   |
| 0.1       | 2               | 20               | 98.37  |
| 0.014     | 3               | 18               | 98.1   |
| 0.0128    | 4               | 16               | 97.83  |
| 0.0125    | 6               | 14               | 97.29  |
| 0.0122    | 7               | 14               | 97     |
| 0.0115    | 8               | 11               | 96.75  |
| 0.0112    | 9               | 9                | 96.48  |
| 0.0104    | 10              | 7                | 96.21  |
| 0.0101    | 11              | 5                | 95.94  |
| 0.009     | 12              | 2                | 95.67  |
| 0.008     | 17              | 0                | 94.32  |

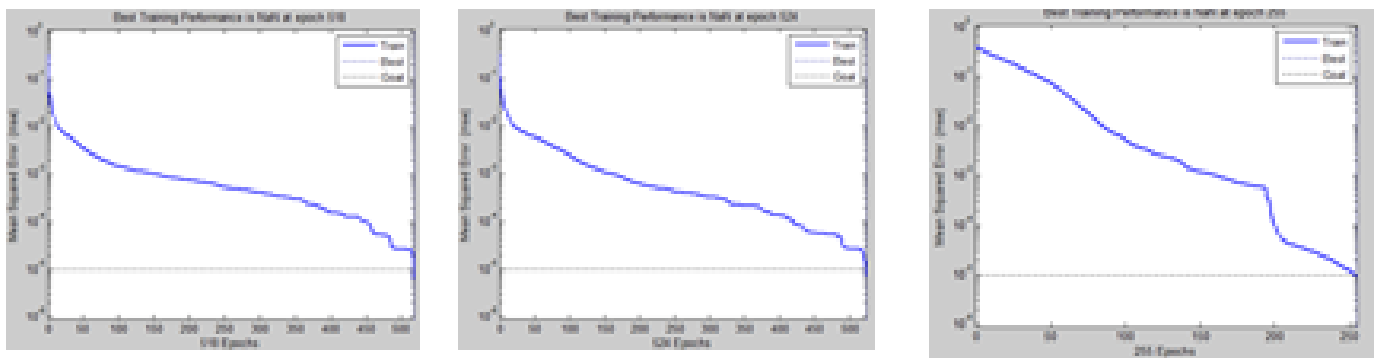


Figure 13. Performance analysis of fbpnn according to number of neurons.

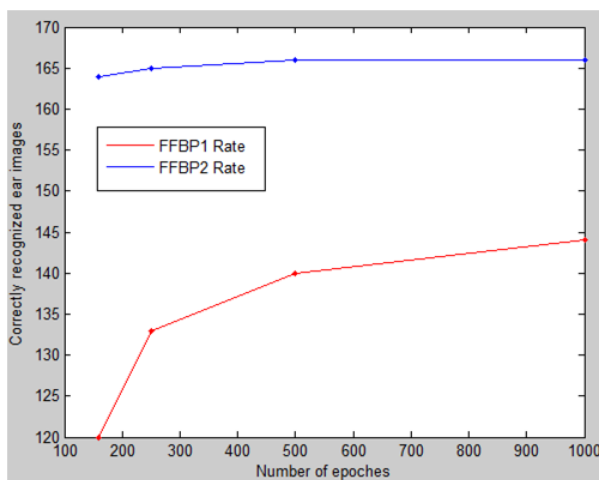


Figure 14. Comparative of different proposed FFBP networks' rate.

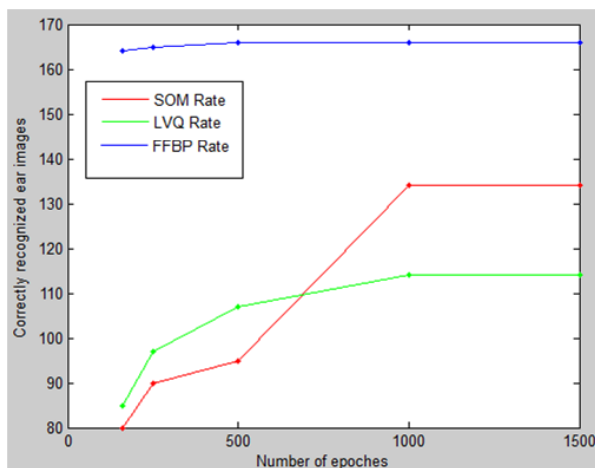


Figure 15. Comparative of different proposed neural networks' rate.

FFBPNN of the three cases (30th,90th,150th) of feature vector. The result are shown in Figures 14,15 and 16 respectively.

If we set the threshold value on 0.01, then the number of false rejected samples on chart MSE1 will be 8. However this number will be 1 if the threshold becomes 0.021.

In MSE2 chart; if the threshold is 0.01, then the number of false rejected samples is 5, and becomes 1 if the threshold is 0.02.

In MSE3 chart; if the threshold value is 0.05, then the number of

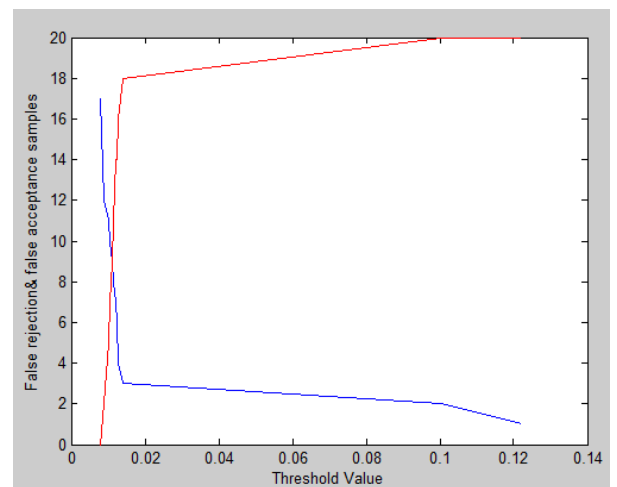


Figure 16. False acceptance (red) and False rejection (blue).

Table 4. Experimental results of different number of features.

| No. features | Training Time | NO. true samples | Rate % |
|--------------|---------------|------------------|--------|
| 30           | 2:16          | 316              | 85.4   |
| 90           | 2:27          | 326              | 88.1   |
| 150          | 5:23          | 347              | 93.7   |
| All (300)    | 3:13          | 366              | 98.9   |

Table 5: Experimental results of different classifier performance.

| Net Type | epochs | neurons    | Training Time | Rate % |
|----------|--------|------------|---------------|--------|
| FFBP     | 1000   | [2000 200] | 5:14          | 95.8   |
| FFBP     | 524    | [1000 200] | 3:13          | 98.9   |
| FFBP     | 255    | [500 10]   | 0:36          | 98.6   |
| SOM      | 700    | 64         | 8:22          | 78.4   |
| SOM      | 1000   | 64         | 12:03         | 81.6   |
| SOM      | 1000   | 144        | 17:34         | 88.9   |
| LVQ      | 500    | 168        | 32:13:00      | 81.6   |
| LVQ      | 800    | 196        | 44:11:00      | 81.1   |
| LVQ      | 1000   | 144        | 44:32:00      | 83.5   |
| K-NN     | -      | -          | -             | 90     |

rejected samples is 2, while this number becomes

1 if the threshold set on 0.1.

### Determine the best classifier

Selecting the best classifier depends on many things such as recognition rate, training time, classifier parameters etc. In the

**Table 6:** Comparative study.

| Researcher                                  | Database Size                            | Used Approach                                                                              | Weak points                                                                                                                                               | Rate                                  |
|---------------------------------------------|------------------------------------------|--------------------------------------------------------------------------------------------|-----------------------------------------------------------------------------------------------------------------------------------------------------------|---------------------------------------|
| A. Ianarelli [10], 1989                     | 10 000 images                            | Manual Ear measurement                                                                     | Very limited and manual                                                                                                                                   | -                                     |
| Burge and Burger [1], 1998                  | -                                        | Neighborhood graph from Voronoi diagrams                                                   | the edges detected from ear image can be very different even in presence of relatively small changes in camera-to-ear orientation or lighting             | Not mentioned                         |
| Moreno [13], 1999                           | 168 images (6 images for 28 individuals) | Compression Networks(Neural Network)                                                       | Does not Take in account illumination and pose variations Does not Take in account occlusion by hair or earning                                           | 93%                                   |
| Yuizono [18], 2002                          | 660 images for 110 individuals           | Genetic search                                                                             | Does not Take in account illumination and pose variations Does not Take in account occlusion by hair or earning                                           | 100%                                  |
| Victor [16], 2002                           | 76 images                                | PCA                                                                                        | Does not Take in account occlusion by hair or earning                                                                                                     | Ear 40%<br>Face 80%                   |
| Chang [3], 2003                             | 197-image training sets                  | PCA “EIGEN-FACES” AND “EIGEN-EARS”                                                         | Does not Take in account occlusion by hair or earning                                                                                                     | Ear 71.6%<br>Face 70.5%<br>Both 90.9% |
| Choras [5], 2005                            | 240 images                               | geometrical approaches                                                                     | Does not Take in account occlusion by hair or earning                                                                                                     | 100%                                  |
| S. M. S. Islam [11], 2008                   | 200 images from 100 individuals          | AdaBoost, Algorithm, Iterative Closest Point (ICP)                                         | Although system takes in account illumination and pose variations, The ICP algorithm can't deal with big pose variations. Depending on prepared database. | 93%                                   |
| S. A. Daramola, O. D. Oluwaninyo, [7], 2011 | 1050 images from 350 individuals         | Haar wavelet transform, Back propagation neural network                                    | Does not Take in account illumination and pose variations Does not Take in account occlusion by hair or earning                                           | 98%                                   |
| Kumar, [12] 2012                            | 753 images                               | Phase encoding with Log Gabor lters                                                        | -                                                                                                                                                         | 95.93%                                |
| Our work                                    | 445 images                               | Likelihood skin detector and morphological operation, Hybrid PCA-Wavelets Features, FFBPNN | (for reviewer)                                                                                                                                            | 98.90%                                |

following, a simple comparative is introduced between different classifiers used in our system in order to specify the best classification method (Table 5).

It can be said that the best classification method is the feed forwarding back propagation neural networks. SOM has a good clustering effectiveness due to its competitive single layer where the similar neurons clustered together, but it still require long training time. LVQ networks have a linear layer (as output) in addition to the competitive one, this will look good for our system, but the result shows that LVQ gives low rates. Nearest neighbor classifier achieved very good performance and can be used instead of neural network to avoid the long training time. However, the FFBPNN has better rate than K-nn.

Regarding to the neural networks, Figure17 illustrates how performance differs using three different networks which are the FFBPNN, SOM and LVQ.

### Comparative study

We compare our study with similar ones in the area of ear recognition.

Table 6 illustrates the main differences between our system and recently ear recognition ones.

### Conclusion

In this paper, a full human recognition based ear images was introduced. The ear images were obtained by means of likelihood skin detector and morphological operations. The features were extracted using PCA-wavelet algorithm, then a FFBPNN was trained by these features. The experiments were done on 370 test images and the system achieved 98.9% recognition rate.

### References

1. Burge M, Burger W (1998) Ear biometrics, In Biometrics: Personal Identification in Networked Society. In: Jain A (Ed's.) Kluwer Academic Publishers.
2. Bustard JD, Nixon MS (2010) Toward Unconstrained Ear Recognition from Two-

Dimensional Images, Ieee Transactions On Systems, Man, and Cybernetics-Part A: Systems and Humans 40: 486-494.

3. Chang K, Bowyer KW, Sarkar S, Victor B (2003) Comparison and Combination of Ear and Face Machine Image in Appearance-Based Biometrics', IEEE Transaction on pattern Analysis and machine Intelligence 25: 1160-1165.
4. Chi-Fa Chen (2003) Combination of PCA and Wavelet Transforms for Face Recognition on 2.5D Images', Palmerston North: 343-347.
5. Choras M (2005) Ear Biometrics Based on Geometrical Feature Extraction “, Electronic Letters on Computer Vision and Image Analysis. pp: 84-95.
6. Conaire CO, O'Connor NE, Smeaton AF (2007) Detector adaptation by maximising agreement between independent data sources, In CVPR. IEEE Computer Society.
7. Daramola SA, Oluwaninyo OD (2011) Automatic Ear Recognition System using Back Propagation Neural Network, *International Journal of Video & Image Processing and Network Security* IJVIPNS-IJENS, 11: 28-32.
8. Elgammal A, Muang C, Hu D (2009) Skin Detection - a Short Tutorial, Encyclopedia of Biometrics by Springer-Verlag Berlin Heidelberg.
9. Hurley DJ, Nixon MS, Carter JN (2005) Force field feature extraction for ear biometrics, Computer Vision and Image Understanding 98: 491-512.
10. Iannarelli A (1989) Ear Identification, Forensic Identification Series, Paramount Publishing Company, Fremont, California.
11. Islam SMS, Bennamoun M, Mian AS, Davies R (2008) A Fully Automatic Approach for Human Recognition from Profile Images Using 2D and 3D Ear Data”, Proceedings of 3DPVT'08 - the Fourth International Symposium on 3D Data Processing 20: 131-135.
12. Kumar A, Wu C (2012) Automated human identification using ear imaging. Pattern Recognition: 956-968.
13. Moreno B, Sanchez A, Velez JF (1999) On the Use of Outer Ear Images for Personal Identification in Security Applications”, IEEE 33rd Annual International Carnahan Conference on Security Technology: 469-476.
14. Nanni L, Brahnam S (2000) A Genetic Algorithm for Creating a Set of Color Spaces for Ear Authentication”, Int'l Conf. IP, Comp. Vision, and Pattern Recognition IPCV'09 200: 971-975.
15. Rahman M, Islam R, Bhuiyan NI, Ahmed B, Islam A (2007) Person Identification Using Ear Biometrics, *International Journal of the Computer, the Internet and Management* 15: 1-8.
16. Victor B, Bowyer KW, Sarkar S (2002) An Evaluation of Face and Ear Biometrics Proc. Int'l Conf. Pattern Recognition: 429-432.

17. Xie Z, Mu Z (2008) Ear Recognition Using LLE and IDLLE Algorithm", ICIC LNCS 1: 460-465.
18. Yuizono T, Wang Y, Satoh K, Nakayama S (2002) Study on Individual Recognition for Ear Images by using Genetic Local Search", In Proc. Of the 2002 Congress on Evolutionary Computation: 237-242.