

Facial Expression Recognition Using FAPs-Based 3DMMM

Hamimah Ujir and Michael Spann

Abstract A 3D modular morphable model (3DMMM) is introduced to deal with facial expression recognition. The 3D Morphable Model (3DMM) contains 3D shape and 2D texture information of faces extracted using conventional Principal Component Analysis (PCA). In this work, modular PCA approach is used. A face is divided into six modules according to different facial features which are categorized based on Facial Animation Parameters (FAP). Each region will be treated separately in the PCA analysis. Our work is about recognizing the six basic facial expressions, provided that the properties of a facial expression are satisfied. Given a 2D image of a subject with facial expression, a matched 3D model for the image is found by fitting them to our 3D MMM. The fitting is done according to the modules; it will be in order of the importance modules in facial expression recognition (FER). Each module is assigned a weighting factor based on their position in priority list. The modules are combined and we can recognize the facial expression by measuring the similarity (mean square error) between input image and the reconstructed 3D face model.

Keywords Modular PCA • 3D Morphable model • 3D Facial expression recognition

H. Ujir (✉)

Universiti Malaysia Sarawak, Kota Samarahan, Sarawak, Malaysia
e-mail: uhamimah@fit.unimas.my

M. Spann

University of Birmingham, Birmingham, UK
e-mail: m.spann@bham.ac.uk

1 Introduction

Facial expression recognition (FER) deals with the application of facial motion and facial feature deformation into abstract classes that are purely based on visual information [1]. Facial expression studies are beneficial to various applications. Among the applications are physiological studies, face image compression, synthetic face animation, robotics as well as virtual reality.

In this chapter, we propose a novel approach for FER called 3D Modular Morphable Model (3DMMM) that combines three advances in the face processing field: 3D Morphable Model (3DMM), Modular Principle Component Analysis (MPCA) and Facial Animation Parameters (FAP).

There have been numerous works in this area. However, our work has the following differences with others: (1) the fitting of 3D shape is done according to the modules and therefore each module has their own eigenvalues and eigenvectors; and (2) each module are given priority in the fitting process which depends on its importance in recognizing facial expression.

The outline of the chapter is as follows: In the second section, the three separate advances are discussed. The framework of this work is explained by the flowchart can be found in the third section. In the fourth section is about the database description. The experiment as well as its analysis is presented in fifth section. Finally, we give the conclusion as well as the limitations and future works.

2 Related Works

2.1 Modular PCA

Using PCA, a face is represented by a linear combination of physical face geometries, S_{mode} and texture images, T_{mode} and both models are within a few standard deviations from their means. PCA is indeed a promising approach in for face analysis. It is fast, reliable and able to produce good results. However, according to Mao et al. [2], PCA does not cope well with variations of expression, facial hair, and occlusion. Thus, we chose a slightly different PCA version which is MPCA to cover all variations of six basic facial expressions.

A similar concept used in this chapter can be found in Tena et al. [3] where they also used a collection of PCA sub-models that are independently trained but share boundaries. The segmentation of the face is a data-driven where the correlation and connection of the vertices are rated based on the range of motion, emotional speech and FAC sequences. The highly correlated and connected vertices form compact regions and compressed by PCA. Their findings strengthen the hypothesis that region-based model is better than holistic approach. However, no FER results recorded as this work is developed for animation purpose.

Zhao et al. [4] state that face recognition using eigenmodules (i.e. mouth, nose and eyes) showed improvement in facial recognition compared to using only eigenfaces. Other work using MPCA showed a significant result especially when there are large variations in facial expression and illumination is by Gottumukkal and Asari [14].

Most work in face processing is based on linear combination approach. Employing PCA on one whole face is like learning a face as one big module. In other word, the local features and its holistic information are not being taken full advantage of.

One major prominent feature of MPCA is it yields new modules to recognize different facial expressions which could be the new addition to the existing facial expression in the training set. The new modules here are the combination of different modules. Besides that, with MPCA, a smaller error is generated compared to conventional PCA as it pays more attention to the local structure [5].

However, according to Gottumukkal and Asari [14], MPCA is known as not giving a significant advancement in pose and orientation problem. It also requires the location of each facial feature to be identified initially. In their work, they also stated that if the face images are divided into very small regions the global information of the face may be lost and the accuracy of this approach is no longer acceptable. Thus, choosing the number of modules to represent a face is also important. MPCA in their work has been employed in face recognition area.

King and Xu [5] divided a face into 4 modules; centre-of-the-left-eye, centre-of-the-right-eye, tip-of-the-nose and centre-of-the mouth feature points. According to them, MPCA generates a smaller error as it pays more attention to the local structure. Chiang et al. [6] used 5 modules which include the left eye, the right eye, the nose, the mouth, and the bare face with each facial module identified by a landmark at the module centre. In this work, a face is divided into several modules according to different facial features which are categorized based on the Facial Animation Parameters (FAPs).

2.2 Facial Animation Parameters

A facial expression is about the deformation of facial features and muscles. Different combination of facial features and muscles produces different type of facial expressions. For instance, how to differentiate between the true smile and polite smile since both types of smile share the same action unit (AU) deformation which is the lip corner puller. In this case, the cheek raiser AU needs to be checked; if it's in action, the true smile is performed and vice versa.

Different works state different AUs are involved in six basic facial expressions which can be found in [Appendix 2](#). Ekman and Friesen [7] introduced the basic AUs involved and in time, other researchers add/deduct certain AUs to represent the facial expressions. We believed this has to do with the intensity of the facial expressions itself, for instance different works might focused on true smile while

others focused on polite smile. However, no report about this matter is found to date. Zhang et al. [8] divided the AUs into two, the primary and auxiliary type. The primary AUs are the AUs that strongly pertinent to one of the facial expressions without ambiguity while the supplementary cues are the auxiliary type.

FAPs are a set of parameters, used in animating MPEG-4 models that define the reproduction emotions, expressions and speech pronunciation. It gives the measurement of muscular action relevant to the AU and it provides temporal information that is needed in order to have a life-like facial expression effect [8]. Figure 1 shows FAPs and feature points on a neutral face and both define Facial Animation Parameters Unit (FAPU).

FAPs represent a complete set of basic facial actions, such as stretch nose, open or close eyelids, and therefore allow the representation of most natural facial expressions. In order to relate the FAPs and AUs, a mapping between both can be found in Appendix 1. The red coloured numbers in Appendix 1 denote that the AUs which are present for that expression in all previous works. Since different works defined different AUs for the basic facial expressions, hence it is affecting the decision on which FAPs to be monitored. In this work, we followed Ekman and Friesen’s [7] work where only 13 AUs were considered.

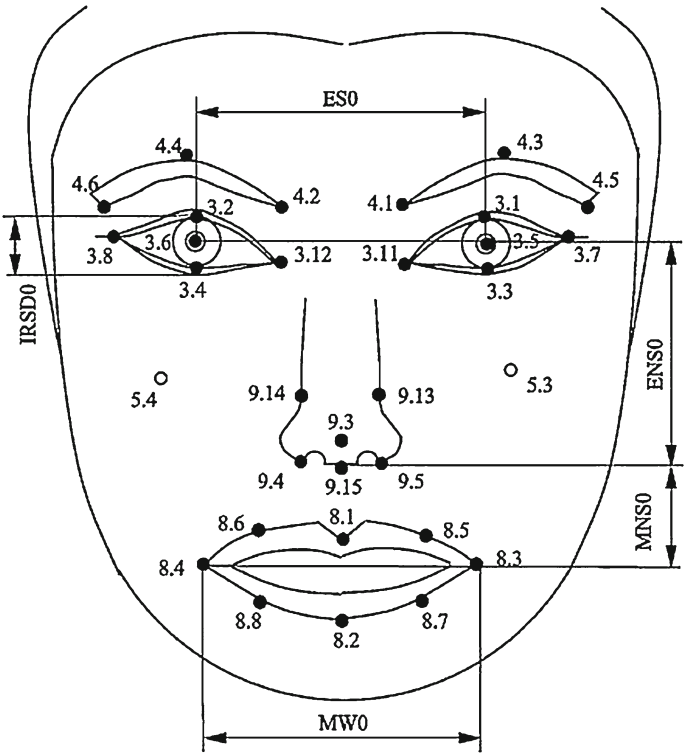


Fig. 1 A neutral face with FAPs and feature points to define FAPU [15]

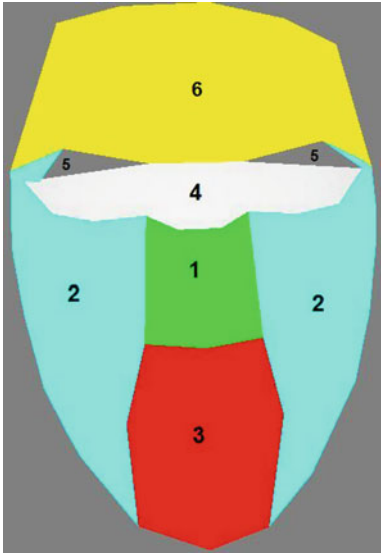


Fig. 2 Facial modules of a face

Table 1 Face modules based on FAPs group

Module	Facial Animation Parameters (FAP)
1	stretch_l_nose, stretch_r_nose
2	lift_l_cheek, lift_r_cheek
3	open_jaw, raise_b_midlip, stretch_r_cornerlip, raise_l_cornerlip, raise_r_cornerlip, push_b_lip, stretch_l_cornerlip, depress_chin, raise_b_midlip_o, stretch_r_cornerlip_o, raise_l_cornerlip_o, raise_r_cornerlip_o, stretch_l_cornerlip_o
4	close_t_r_eyelid, close_t_l_eyelid, close_b_r_eyelid, close_b_l_eyelid
5	raise_r_i_eyebrow, raise_r_o_eyebrow, squeeze_r_eyebrow, raise_l_i_eyebrow, raise_l_o_eyebrow, squeeze_l_eyebrow

Figure 2 shows different colours for every module and each colour defined the modules priority. In this work, six FAP-based modules are used to represent a whole face, see Table 1. Each module contains the facial features that correspond to the FAP in that module. Forehead module is not mentioned in Table 1 as no FAPs involved in six facial basic emotions come from this region. However we also include this module in this work because the end outcome needs to have a complete face as well as to see how it influences each of the expression. In group one, the two FAPs only exists in disgust and anger expression and when there is a change for these two parameters, it means the subject is showing the disgust or anger expression.

2.3 FAPs-Based 3D MMM

3D MM is based on a series of 3D scans example represented in an object centred system and registered to a single reference [9]. It will be used for MPCA-based representation of faces; this combination of algorithms will then be named as the 3D Modular Morphable Model (3D MMM).

In this work, a Modular PCA (MPCA) is implemented where each face module is treated separately in the PCA process. The face modules are decided based on the facial features which are categorized according to the Facial Animation Parameters (FAP). Basic related concepts of FAP are described in Lavagetto and Pockaj [10].

The training of MPCA is similar to conventional PCA with the algorithm applied to each of the six groups in Table 1. For convenience, six disjoint sets of facial features are denoted as $P \in \{\text{forehead}, \text{eyebrows}, \text{eyes}, \text{mouth}, \text{cheeks}, \text{nose}\}$. The training examples are stored in terms of x, y, z -coordinates of all vertices in the same modules of a 3D mesh. The following is the example of the nose module:

$$S_{\text{nose}} = \{x_1, y_1, z_1, x_2, y_2, z_2, \dots, x_n, y_n, z_n\} \quad (1)$$

$$S_{\text{nose}} = \bar{S}_{\text{nose}} + \sum_{i=1}^n a_i s_i \quad (2)$$

The linear space of face geometries is denoted in Eq. 2 and it assumes a uniform distribution of the shapes. a_{nose} is the coefficients that determine the variation between 3D nose modules for all faces in the training set.

In this work, our aim is to experiment the decomposition of 3D face shape. Therefore, for appearance part, one modular approach is implemented. Accordingly, the texture vectors are formed from the red, green and blue of all vertices and b (in Eq. 4) is the appearance coefficient that determine the appearance variation for all faces in the training set.

$$T = \{r_1, g_1, b_1, r_2, g_2, b_2, \dots, r_n, g_n, b_n\} \quad (3)$$

$$T = \bar{T} + \sum_{i=1}^n b_i t_i \quad (4)$$

3 Our Framework

PCA algorithms produce a set of values of uncorrelated variables called principal module (PC). All PCs are then ordered so that the first few retain most of the variation in all of original variables while the rest of the modules contain the remaining original variables after all correlation with the preceding PCs has been

subtracted out. The number of PCs is normally chosen to explain at least 90 % of the variation in the training set.

The fitting process of a new image to the model involved projecting the 2D image onto the subspace (which is called the “face space”) and then finding the minimum of the distance all of the faces stored in the database and the closest matching one is recognised. With the number of subjects involved, as well as large variation of expressions, poses and illumination in the training set, the number of PCs to be considered is rather high as it needs to cover all variation of faces in the dataset. Rationally, the linear computational cost is linear with the number of PCs. With MPCA, the number of PCs to be considered for each module is lower when compared to calculating PCs for a whole face.

Figure 3 describes our 3D Modular Morphable Model Framework. For 3D shape data in training set, they were all decomposed into six modules then each module will go through PCA calculation. However, the appearance data will go straight to PCA process without decomposition process.

Given a 2D image of a subject with facial expression, a matched 3D model for the image is found by fitting them to our 3D MMM. Again, the 3D shape of a probe image is decomposed into six modules, then the matching shape for each modules are found by projecting each modules into their respected face modules space.

As mentioned before, the 3D shape fitting is done according to the modules; it will be in order of the importance modules in FER. Each module is assigned a weighting factor based on their position in priority list. For instance, three modules are involved to decide the new 3D facial landmarks for eyes module which are the eyebrows, cheeks and nose. Nose will be given a higher weighting factor because of its position in priority list and then followed by cheeks and eyebrows. While for appearance, the fitting is done for a whole face, similar to a conventional PCA fitting. Finally, a 3D shape and a texture are combined and a new 3D face model is generated.

4 Database Description

A multi-attribute database developed by Savran et al. [18] from Bogazici University, Turkey called The Bosphorus Database is used. The data is acquired using Inspeck Mega Capturor II 3D which is a commercial structured-light based 3D digitizer device. The Bosphorus Database contains 24 facial landmark points; provided that they are visible in the given scan (i.e., the right and left ear lobe cannot be seen from frontal pose). It provides a rich set of expressions, systematic variation of poses and different types of realistic occlusions. Each scan is intended to cover one pose and/or one expression type. Thirty-four facial expressions are composed of a wisely chosen subset of Facial Action Units (FAU) as well as the six basic emotions. Besides the facial expression data, this database contains different occlusion and head poses data.

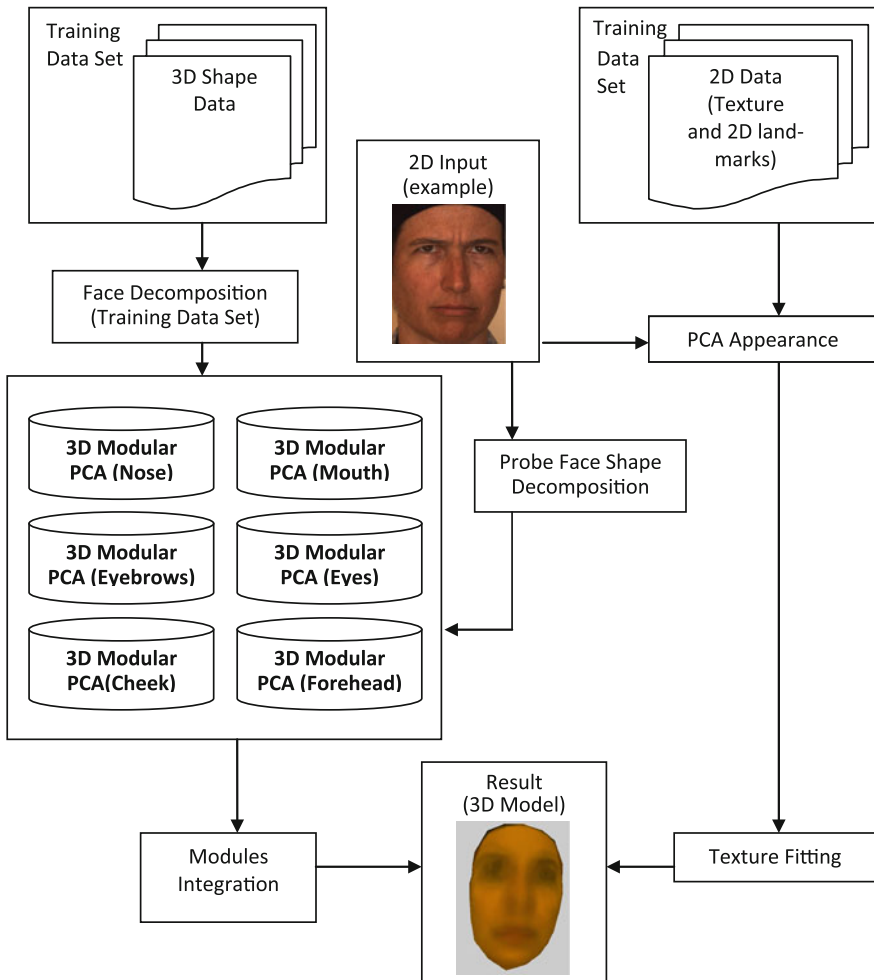


Fig. 3 3D modular morphable model framework

In this work, our 3DMMM is derived from statistics computed on 54 subjects with 6 different expressions which are anger, disgust, fear, happy, sad and surprise. Figure 4 shows a subject with six different facial expressions. Extra work need to be done as we decided to add 6 more facial landmarks which are the centre of pupils, the highest and lowest point on both eyes. These 6 extra facial landmarks are needed as it involve in the selected FAPs list.

We are dealing with two types of information, 2D and 3D data, and both data need to go under a few processes before they are ready to be used in modular PCA computation as well as the fitting procedure. The 3D points of every faces are aligned to each other and 115 of 3D points and 210 meshes are used to represent one whole face. The aligned 3D points are then divided into six modules. For each

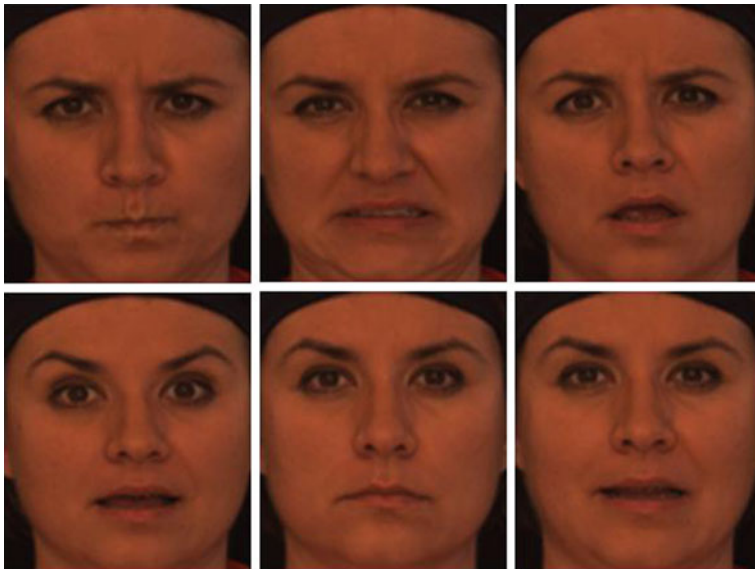


Fig. 4 Six different emotions of a subject. *Top row* anger, disgust, fear. *Bottom row* surprise, neutral and few combinations of deformation

module, the eigenvalues and its corresponding eigenvectors are computed. Out of certain amount of eigenvalues, only 97 % of the eigenvalues of the whole training data are kept and used in the next step which is fitting.

In two dimensional data, two important data are needed: (1) 2D-points that mark the facial features on a texture and (2) texture (RGB values) with a size of 50×50 pixels are used. The 2D points are warped using Thin Plate Spline (TPS) algorithm to ensure the colour profiles are obtained cross a free shape patch. All training texture are also normalized, this way all training texture have the same brightness and contrast. The RGB values of the texture are extracted from the normalized texture and eigenvalues and its corresponding eigenvectors are computed.

We can recognize the facial expression by measuring the similarity (mean square error) between input image and the reconstructed 3D face model. Following Blanz et al. [11] works', rigid transformation and perspective projective are implemented in the fitting process. A further discussion of fitting a new 2D image can be found in Blanz et al. [11].

5 Results and Analysis

To assess the viability of this approach, we have performed the experiments to recognize six facial expressions. In the test set, there are 4 subjects and each subject shown 6 facial expressions. Our FER experiments are carried out in a

Table 2 Facial expression recognition rate (%) using MPCA

	Anger	Disgust	Fear	Happy	Sad	Surprise
<i>Nose</i>						
NNS	75	50	25	50	100	25
<i>Mouth + Jaws</i>						
NNS	25	50	50	50	50	75
<i>Eyes</i>						
NNS	50	25	50	75	25	25
<i>Eyebrows</i>						
NNS	25	50	25	25	50	50
<i>Cheeks</i>						
NNS	25	50	50	75	50	25
<i>Forehead</i>						
NNS	25	25	25	50	50	100

person-independent manner as all test subjects were not in the training set. According to Pantic and Rothrantz [17], person independent experiment in FER is more challenging than person-dependent approach. There are no rejections, only correct recognition or false recognition as no threshold is used in this experiment.

Several experiments have been carried out and each of the results can be find in the table below: (1) FER using MPCA (2) FER for using only PCA and (3) FER after integrating all modules. Table 2 shows the FER rate for each of the face modules using nearest neighbour search approach which is calculated in Euclidean distance measure. In terms of specific expressions, recognition rate of surprise is the best and anger and fear are the worst. Nose module gives the best result and the worst is eyebrows module. The eyebrows module is the smallest region on face. We believed that the global information is lost because of the small region size and therefore affecting the FER rate.

Table 3 shows the FER rate for non-modular PCA while Table 4 shows the FER rate for combined modules. We can see that the latter one improves the FER rate compared to the former one. However, the FER rate for surprise is similar for both approaches. Although each module give a quite promising output (refer Table 2) but when they were all combined, the result is not sustained. We believed the assigned weighting factor has affected each module in the fitting process.

Table 3 Facial expression recognition rate (%) for non-modular PCA

	Anger	Disgust	Fear	Happy	Sad	Surprise
NNS	25	25	25	50	25	50

Table 4 Facial expression recognition rate (%) for combined modules

	Anger	Disgust	Fear	Happy	Sad	Surprise
NNS	50	50	50	100	50	50



Fig. 5 Input image with anger expression (*left*). 3D model with anger appearance (*right*)

Figure 5 shows a test image with anger expression (left) and the generated 3D model (right) while Fig. 6 shows the fear expression. Another noticeable outcome from the experiment is the wrinkle and dimple feature. In real life, these two features are the two keys that help in identifying people and it also become one of the important feature that project certain facial expression. For example, in Fig. 6, lines of wrinkles on the forehead and a line from inner eyes to the outer cheek are part of the fear expression. However, this cannot be seen in the 3D model generated.



Fig. 6 Input image with fear expression (*left*). 3D model with fear appearance (*right*)

6 Conclusion

This paper explores the potential of facial expression recognition using modular approach. A face is divided into six modules and each module will have their own eigenvalues as well as eigenvector. A test image is divided into the six modules. Each module is assigned a weighting factor based on their position in priority list. The weighting factor is used to integrate the modules and then give us a new face. The system developed also yields various facial expressions even though that certain expression is not in the training set.

There are some limitations in the current work:

- (1) The combined modules only perform more than average but still better than non-modular approach. We believe using only 3D facial landmarks to measure facial expression is just not enough to capture the facial expression information.
- (2) The weighting factors assigned give a rather high impact on each module when combined. For instance, though the eyebrows module is put on the second last position in the priority list, it does affect the facial landmarks in forehead and eyes modules in the fitting process.
- (3) Only 4 subjects with 6 facial expressions are tested in this work due to the limited data.
- (4) No appearance features like wrinkle and dimple is generated in this work.
- (5) The texture computation component in this work is rather time-consuming compared to the shape. This is due non-modular PCA approach used in texture component as the number of PCs that need to be calculated is quite large. Our future work will emphasize in finding a purely 3D shape feature to be used in FER as the 3D facial landmarks are not enough to measure facial expression. In order to have an effective classification system, the modules fitting process needs to be improved and number of test image will be added. We will pursue these three aspects in the future.

Appendix 1

AU	Description	FAP number	FAP name	Module
1	Inner brow raiser	31	raise_l_i_eyebrow	5
		32	raise_r_i_eyebrow	
2	Outbrow raiser	35	raise_l_o_eyebrow	5
		36	raise_l_o_eyebrow	
4	Brow lower	31_	raise_l_i_eyebrow	5
		32_	raise_r_i_eyebrow	
		37	squeeze_l_eyebrow	
		38	squeeze_r_eyebrow	

(continued)

(continued)

AU	Description	FAP number	FAP name	Module
5	Upper lid raiser	19_	open_t_l_eyelid (close_t_l_eyelid)	4
		20_	open_t_r_eyelid (open_t_r_eyelid)	
6	Cheek raiser	19	close_t_l_eyelid	5
		20	close_t_r_eyelid	
		41	lift_l_cheek	
		42	lift_r_cheek	
7	Lid tighter	21	close_b_l_eyelid	4
		22	close_b_r_eyelid	
9	Nose wrinkler	61	stretch_l_nose	1
		62	stretch_r_nose	
10	Upper lip raiser	59	raise_l_cornerlip_o	3
		60	raise_r_cornerlip_o	
12	Lip corner puller	59	raise_l_cornerlip_o	3
		60	raise_r_cornerlip_o	
		53	stretch_l_cornerlip_o	
		54	stretch_r_cornerlip_o	
15	Lip corner depressor	59_	lower_l_cornerlip (raise_l_cornerlip_o)	3
		60_	lower_r_cornerlip (raise_r_cornerlip_o)	
16	Lower lip depressor	5	raise_b_midlip	3
		16	push_b_lip	
17	Chin raiser	18	depress_chin	3
20	Lip stretcher	53	stretch_l_cornerlip	3
		54	stretch_r_cornerlip	
		5	raise_b_midlip	
23	Lip tighter	53_	tight_l_cornerlip	3
		54_	tight_r_cornerlip	
24	Lip pressor	4	lower_t_midlip	3
		16	push_b_lip	
		17	push_t_lip	
25	Lip apart	3	open_jaw(slight)	3
		5_	lower_b_midlip(slight)	
26	Jaw drop	3	open_jaw(middle)	3
		5_	lower_b_midlip(middle)	
27	Mouth stretch	3_	open_jaw(large)	3
		5_	lower_b_midlip(large)	

Appendix 2

	Ekman and Friesen [7]	Raouzaoui et al. [12]	Zhang et al. [8]		Deng and Noh [13]	Lucey et al. [16]	Velusamy et al. [19]
			Primary	Auxiliary			
Anger	$4 + 5 + 7 + 23$	$2 + 4 + 5 + 7 + 17$	$2 + 4 + 7 + 23 + 24$	$17 + 25 + 26 + 16$	$2 + 4 + 7 + 9 + 10 + 20 + 26$	$4 + 5 + 15 + 17$	$23 + 7 + 17 + 4 + 2$
Disgust	$9 + 15 + 16$	$5 + 7 + 10 + 25$	$9 + 10$	$17 + 25 + 26$	NIL	$1 + 4 + 15 + 17$	$9 + 7 + 4 + 17 + 6$
Fear	$1 + 2 + 4 + 5 + 20 + 26$	$4 + 5 + 7 + 24 + 26$	$20 + (1 + 5) + (5 + 7) + 12$	$4 + 5 + 7 + 25 + 26$	$1 + 2 + 4 + 5 + 15 + 20 + 26$	$1 + 4 + 7 + 20$	$20 + 4 + 1 + 5 + 7$
Happiness	$6 + 12$	$26 + 12 + 7 + 6 + 20$	$6 + 12$	$16 + 25 + 26$	$1 + 6 + 12 + 14$	$6 + 12 + 25$	$12 + 6 + 26 + 10 + 23$
Sadness	$1 + 4 + 15$	$7 + 5 + 12$	$1 + 15 + 17$	$4 + 7 + 25 + 26$	$1 + 4 + 15 + 23$	$1 + 2 + 4 + 15 + 17$	$15 + 1 + 4 + 17 + 10$
Surprise	$1 + 2 + 5B + 26$	$26 + 5 + 7 + 4 + 2 + 15$	$5 + 26 + 27 + (1 + 2)$	NIL	$1 + 2 + 5 + 15 + 16 + 20 + 26$	$1 + 2 + 5 + 25 + 27$	$27 + 2 + 1 + 5 + 26$

References

1. Fasel B, Luetttin J (2003) Automatic facial expression analysis: a survey. *Pattern Recogn* 36(1):259–275
2. Mao X, Xue Y, Li Z, Huang K, Lv S (2009) Robust facial expression recognition based on RPCA and AdaBoost. In: 10th workshop on image analysis for multimedia interactive services
3. Tena R, De la Torre F, Matthews I (2011) Interactive region-based linear 3D face models. *ACM Trans Graph* 30(4):76
4. Zhao W, Chellappa R, Phillips PJ, Rosenfeld A (2003) Face recognition: a literature survey. *ACM Comput Surv* 35(4):399–458
5. King I, Xu L (1997) Localized principal module analysis learning for face feature extraction and recognition. In: Proceedings of workshop 3D computer vision, p 124
6. Chiang C-C, Chen Z-W, Yang C-N (2009) A module-based face synthesizing method. In: APSIPA annual summit and conference, p 24
7. Ekman P, Friesen W (1978) Facial action coding system: a technique for the measurement of facial movement. Consulting Psychologists Press, Palo Alto
8. Zhang Y, Ji Q, Zhu Z, Yi B (2008) Dynamic facial expression analysis and synthesis with MPEG-4 facial animation parameters. *IEEE Trans. Circuits Syst Video Technol* 18(10): 1383–1396
9. Romdhani S, Pierrard J-S, Vetter T (2005) 3D morphable face model, a unified approach for analysis and synthesis of images. In: Wenyi Zhao RC (ed) *Face processing: advanced modeling and methods*. Elsevier
10. Lavagetto F, Pockaj R (1999) The facial animation engine: towards a high-level interface for the design of MPEG-4 compliant animated faces. *IEEE Trans Circuits Syst Video Technol* 9(2):277–289
11. Blanz V, Scherbaum K, Seidel H (2007) Fitting a morphable model to 3D scans of faces. *Comput Vis IEEE Int Conf* 0:1–8
12. Raouzaoui A, Tsapatoulis N, Karpouzis K (2002) Kollias S (2002) Parameterized facial expression synthesis based on MPEG-4. *EURASIP J Appl Sig Proc* 1:1021–1038
13. Deng Z, Noh J (2008) Computer facial animation: a survey. In: Deng Z, Neumann U (eds) *Data driven 3D facial animation*. Springer, pp 1–28
14. Gottumukkal R, Asari VK (2003) An improved face recognition technique based on modular PCA approach. *Pattern Recognit Lett* 25(4):429–436
15. ISO/IEC IS 14496-2 Visual (1999) <http://kazus.ru/nuke/modules/Downloads/pub/144/0/ISO-IEC-14496-2-2001.pdf>. Assessed 13 Feb 2012
16. Lucey P, Cohn JF, Kanade T, Saragih J, Ambadar Z, Matthews I (2002). The extended Cohn-Kanade dataset (CK +): a complete dataset for action unit and emotion-specified expression. In: IEEE computer society conference on computer vision and pattern recognition workshops (CVPRW), pp 94–101
17. Pantic M, Rothkrantz L (2000) Automatic analysis of facial expressions: the state of the art. *IEEE Trans PAMI* 22:1424–1445
18. Savran A, Alyüz N, Dibeklioglu H, Çeliktutan O, Gökberk B, Sankur B, Akarun L (2008) Biometrics and identity management. In: Schouten B, Juul NC, Drygajlo A, Tistarelli M (eds) *Bosphorus database for 3D face analysis*. Springer, Berlin, pp 47–56
19. Velusamy S, Kannan H, Anand B, Sharma A, Navathe B (2011) A method to infer emotions from facial action units. In: IEEE international conference on acoustics, speech and signal processing (ICASSP), pp 2028–2031