

Hallucination Of Face Images Across Multiple Modalities Using Tensors

1. Ms. S.Gayathri

Assistant Professor, Department of ECE, Sri Sai Ram Engineering College Chennai, India gayathri.ece@sairam.edu.in

2. Ms.K.Jeyapiriya

Assistant Professor, Department of ECE, Sri Sai Ram Engineering College Chennai, India
jeyapiriya.ece@sairam.edu.in

3. Mr.J.Prakash

Assistant Professor, Department of ECE, Sri Sai Ram Engineering College Chennai, India
prakash.ece@sairam.edu.in

Article Info

Volume 82

Page Number: 11545 - 11549

Publication Issue:

January-February 2020

Abstract

The existing face resolution techniques produces images with higher resolution from low resolution face images of a single modality. Single modality of facial images may be low resolution images of a fixed expression, fixed pose and a particular illumination. The technique presented here is hallucinating high resolution facial images over various modalities, i.e., with diverse facial expressions, poses and illumination conditions. This work addresses the following issues of the facial images taking various combinations of persons and expressions. 1. For a person with a single expression, is it possible to develop his other different expressions? 2. For a person with a single expression, is it possible to generate the same expression for different persons? 3. Given a facial image of low resolution, is it possible to recognize the person? To achieve the above issues, the work is concentrated on registering the unregistered raw images using an automatic face alignment algorithm and on developing a generalized tensor. The tensors are decomposed using HOSVD and finally synthesized to high or super resolution by Interpolation. This novel technique not only shows enhanced and superior performance over existing super resolution techniques but shows robustness in coping with hallucination of multi modal images under different practical conditions.

Article History

Article Received: 18 May 2019

Revised: 14 July 2019

Accepted: 22 December 2019

Publication: 21 February 2020

Keywords: HOSVD, Automatic face alignment algorithm, tensors

I. INTRODUCTION

The problems of current facial recognition techniques are that they have undergone various issues because of the following reasons. The quality of the captured image is different from the reference image in the database because of lower capture resolution, different lighting effects, and difference in the profile of the captured image or high noise in the environmental conditions. The older facial recognition systems fail to respond to real time images. A facial image captured by a live surveillance camera from a distance may consist of

a various non linearities because of changes in expression, view point, also called pose and illumination. The absence of high resolution

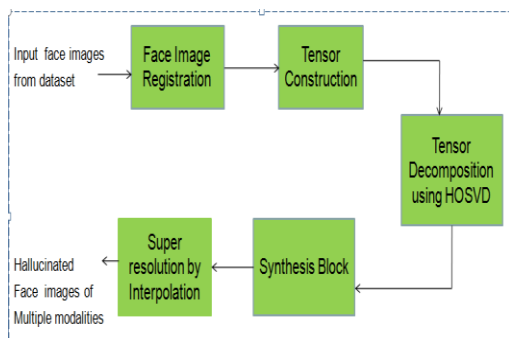
details deteriorates the analysis and recognition of the images. The following issues of the facial images taking various combinations of persons and expressions need to be addressed considering the above situations.

1. For a person with a single expression, is it possible to develop his other different

expressions?

2. For a person with a single expression, is it possible to generate the same expression for different persons?
3. Given a facial image of low resolution, is it possible to recognize the person?

II. THE OVERALL PAPER IS ORGANIZED AS FOLLOWS:



- (i) Register the input face images.
- (ii) Form Tensor A of size $2 \times 3 \times 2160 (54 \times 40)$
- (iii) Decompose the tensor A by flattening to obtain A1, A2 and A3.
- (iv) Apply HOSVD to A1, A2, A3 and obtain $A1 = U1 \times S \times V1^T$

$$A2 = U2 \times S \times V2^T$$

$$A3 = U3 \times S \times V3^T$$

- (v) CONSTRUCT THE CORE TENSOR FROM TENSOR DECOMPOSITION AS FOLLOWS:

$$S = A \times U1^T \times U2^T \times U3^T$$

- (vi) Form person tensor as follows:

$$C = S \times U1 \times U3$$

- (vii) Form expression tensor as follows:

$$B = S \times U1 \times U2$$

- (viii) Obtain high resolution multiple modalities face images using super resolution technique.

A. Block diagram of the system:

This paper is motivated to find solution to the above issues by developing generalized hallucination techniques for facial images across multiple modalities. To achieve this, a unified tensor is developed.

B. Image Registration:

The image registration can be performed in the following steps:

i) Feature extraction – to identify the features needed for further processing like edges, regions etc.,

ii) Feature matching – to check for the correlation of features between the two images.

iii) Mapping function building – to build a mapping function from the above step.

iv) Image registration – to register the image from the mapped one. The image registration problem has two parts, selecting accurate control points and interpolating the image. Area based matching (ABM) and Feature Based Matching (FBM) are the techniques for control point selection. Feature-based matching techniques uses various features of the image obtained from the feature extraction techniques instead of gray values. To register the image, the following steps are to be carried out:

- Compute $F'(u,v)$ and $G'(u,v)$ – To determine the fourier transform of the image $f'(x,y)$ and $g'(x,y)$.
- Compute the product $F'(u,v)$ and $G^{*'}(u,v)$ – Multiply the two images in their frequency domain.
- Compute inverse FFT of the product to obtain cross correlation.
- Locate its peak for image registration.

II. TENSOR – AN OVERVIEW:

A matrix is a framework of $n \times m$ (say, 3×3) numbers surrounded by brackets. A vector is a matrix with one row or column. So there are a bunch of mathematical operations that we can do to any matrix. A matrix is just a 2-D grid of numbers. A tensor can be called as a generalized matrix, a 1-D, a 3-D matrix and even a 0-D matrix (containing only one element), or a higher

dimensional matrices that goes beyond our imagination. A tensor is a mathematical entity that lives in a structure and interacts with other mathematical entities. If one *transforms* the other entities in the structure in a regular way, then the tensor *must obey a related transformation rule*. This “dynamical” property of a tensor is the key that distinguishes it from a normal matrix.

A Tensor of size 2 X 3 X 2160 is constructed.

- 2 Persons
- 3 Expressions
- 160 pixels (54 x 40)

Fig 1 Tensor set for facial images



III. MULTIPLE MODALITIES:

- 1: Constant Face Image
- 2: Smiling Face Image
- 3: Angry Face Image
- 4: Screaming Face Image
- 5: Left Side Lighting Face Image
- 6: Right Side Lighting Face Image

Fig 2. Different modalities of a single person



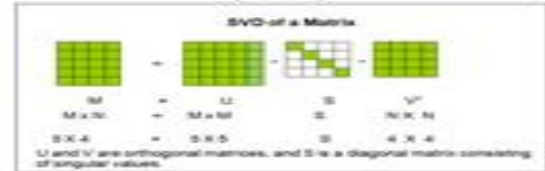
IV. FACIAL EXPRESSION DECOMPOSITION

$U^{(1)}$ is the U^{Person} is the person row vector subspace; $U^{(2)}$ is the $U^{\text{expression}}$ is the expression rowvector subspace ; $U^{(3)}$ is the U^{feature} facial feature row vector subspace ; S is defined as the core tensor and it shows the interrelationship between person, expression and feature subspaces. . Each column vector in each

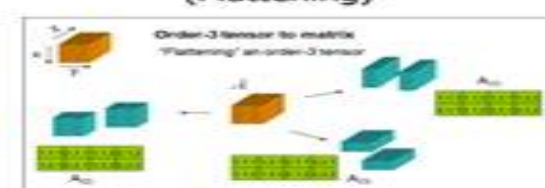
subspace matrix gives the characteristics of the other modalities.

V. SINGULAR VALUE DECOMPOSITION

Singular Value Decomposition (SVD)



Reshaping Tensors (Flattening)



VI. HALLUCINATING FACE IMAGES ACROSS MODALITIES

Local patch- based multi resolution tensor performs the resolution for every face modalities.

Local tensor cannot recover the high frequency informations. A non parametric local patch updating is added by learning the training data set

Suppose that the images with higher resolution that need to be hallucinated for different modalities, are their low-resolution facial images, and any face image with poor-resolution maximum *a posteriori* (MAP) estimation is done for high resolution images.

Every high resolution facial images are reconstructed from the synthesized low resolution images for single modality

A three step sequential solution is herewith achieved wherein the first step is synthesizing the low resolution S1 and S2. Thus given a low

resolution constant expression image, the other modalities, for example different expressions can be generated by hallucinating the facial images using tensors.

Thus given a low resolution constant expression image, the other modalities, for example different expressions can be generated by hallucinating the facial images using tensors. This is accomplished for the training set

VII. FACIAL EXPRESSION HALLUCINATION

The matrix that contains all the facial features with all expressions forms the expression tensor. The person tensor defines the matrices that contains all the facial features.

The tensor is flattened with respect to each dimension and the flattened matrices A1, A2 and A3 are obtained.

i) i) Apply HOSVD to A1, A2, A3 matrices as follows:

ii) $A1 = U1 \times S \times V1T$

$A2 = U2 \times S \times V2T$

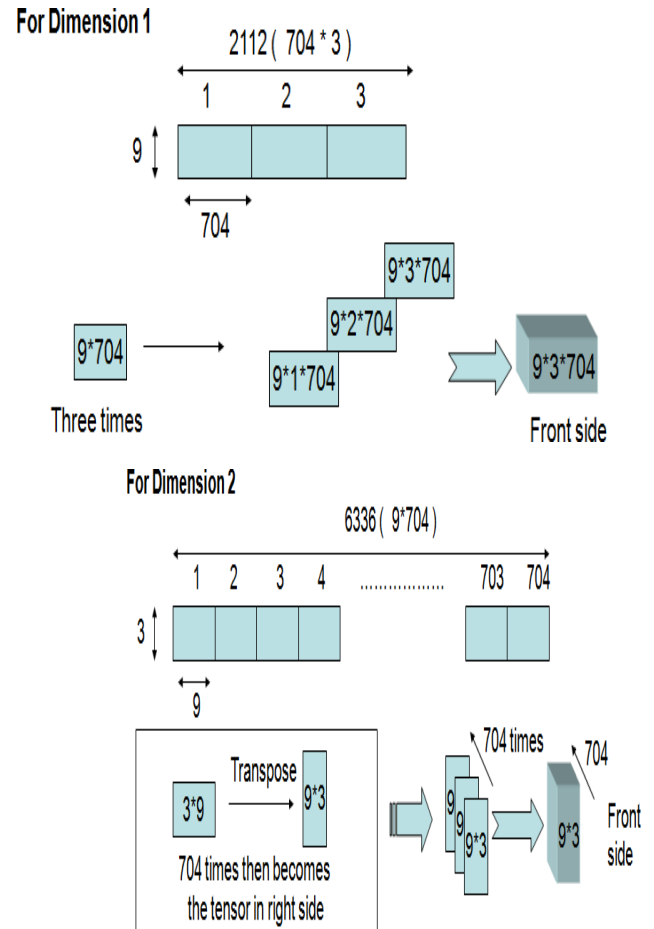
$A3 = U3 \times S \times V3T$

ii) Construct the core tensor from tensor decomposition as follows:

$$S = A \times U1T \times U2T \times U3T$$

The following figures illustrate the construction of tensors considering the different modalities

The flattened matrix are once again converted back to unflattened tensors with respect to the three dimensions



VIII. SUMMARY

In this paper, a method for hallucinating facial images across multiple modalities is proposed. High- and low-resolution training face images multiple modalities are also designed. When a low-resolution face input of a single modality is given, multiple face images poor in resolution of different modalities are done initially using a trained global tensor. Based on these low-resolution face images that are been synthesized, a trained local tensor is used to construct corresponding image of better resolution for each facial modality. For the highest frequency information, high-frequency component residue is further added using nonparametric patch learning from high-resolution training set. The registration process is done by iteratively optimizing affine transformation

parameters, which warp any probing face in raw images to its projection in a PCA subspace. Experimental results show superiority in performance in resolution of facial images both for single and across various modalities of face images.

IX. REFERENCES

- [1] Image resolution enhancement by using discrete and stationary wavelet decomposition H Demirel, G Anbarjafari - IEEE transactions on image ..., 2010 - ieeexplore.ieee.org
- [2] A novel sparse representation based framework for face image super-resolution G Gao, J Yang -
- [3] Neurocomputing, 2014 - Elsevier
- [4] Face recognition by metropolitan police super-recognisers DJ Robertson, E Noyes, AJ Dowsett, R Jenkins... - Plus one, 2016 - journals.plos.org
- [5] Evaluation of image resolution and super-resolution on face
- [6] recognition performance C Fookes, F Lin, V Chandran, S Sridharan - ... Communication and Image ..., 2012 -
- [7] Elsevier
- [8] Face image super-resolution using 2D CCA L An, B Bhanu - Signal Processing, 2014 - Elsevier
- [9] Super-resolution track density imaging of glioblastoma: histopathologic correlation RF Barajas, CP Hess, JJ Phillips... - American Journal ..., 2013 - Am Soc Neuroradiolog
- [10] Hallucinating faces: LPH super-resolution and neighbor reconstruction for residue compensation Y Zhuang, J Zhang, F Wu - Pattern Recognition, 2007 - Elsevier
- [11] Super-resolution of human face image using canonical correlation analysis H Huang, H He, X Fan, J Zhang -
- [12] Pattern Recognition, 2010 - Elsevier
- [13] Face hallucination based on sparse local-pixel structure
- [14] Y Li, C Cai, G Qiu, KM Lam - Pattern Recognition, 2014 -
- [15] Elsevier
- [16] Method of restoring and reconstructing super-resolution image from low-resolution compressed image IK Lim, MG Kang, SC Park - US Patent App. 10/875,218, 2005 - Google Patents
- [17] Kui Jia, Shaogang Gong. "Generalized Face Super- Resolution", IEEE Transactions on Image Processing, 2008.