



JOURNAL OF IMAGE AND GRAPHICS

主办: 中国科学院空天信息创新研究院
中国图象图形学学会
北京应用物理与计算数学研究所

中国图象学报 中国图形学报

2022
07
VOL.27

ISSN1006-8961
CN11-3758/TB



舰船目标检测 P2078



第27卷第7期（总第315期）
2022年7月16日

中国精品科技期刊
中国国际影响力优秀学术期刊
中国科技核心期刊
中文核心期刊

版权声明

凡向《中国图象图形学报》投稿，均视为同意在本刊网站及CNKI等全文数据库出版，所刊载论文已获得著作权人的授权。本刊所有图片均为非商业目的使用，所有内容，未经许可，不得转载或以其他方式使用。

Copyright

All rights reserved by Journal of Image and Graphics, Institute of Remote Sensing and Digital Earth, CAS. The content (including but not limited text, photo, etc) published in this journal is for non-commercial use.

主管单位 中国科学院
主办单位 中国科学院空天信息创新研究院
中国图象图形学学会
北京应用物理与计算数学研究所

主 编 吴一戎
编辑出版 《中国图象图形学报》编辑出版委员会
通信地址 北京市海淀区北四环西路19号
邮 编 100190
电子信箱 jig@aircas.ac.cn
电 话 010-58887035
网 址 www.cjig.cn

广告发布登记号 京朝工商广登字20170218号
总 发 行 北京报刊发行局
订 购 全国各地邮局
海外发行 中国国际图书贸易集团有限公司
(邮政信箱: 北京399信箱 邮编: 100048)
印刷装订 北京科信印刷有限公司

Journal of Image and Graphics

Title inscription: Song Jian | Monthly, Started in 1996

Superintended by Chinese Academy of Sciences
Sponsored by Aerospace Information Research Institute, CAS
China Society of Image and Graphics
Institute of Applied Physics and Computational Mathematics

Editor-in-Chief Wu Yirong
Editor, Publisher Editorial and Publishing Board of Journal of Image and Graphics
Address No. 19, North 4th Ring Road West, Haidian District, Beijing, P. R. China
Zip code 100190
E-mail jig@aircas.ac.cn
Telephone 010-58887035
Website www.cjig.cn

Distributed by Beijing Bureau for Distribution of Newspapers and Journals
Domestic All Local Post Offices in China
Overseas China International Book Trading Corporation
(P.O.Box 399, Beijing 100048, P.R.China)
Printed by Beijing Kexin Printing Co., Ltd.

CN 11-3758/TB
ISSN 1006-8961
CODEN ZTTXFZ

国外发行代号 M1406
国内邮发代号 82-831
国内定价 60.00元



深度学习背景下视觉显著性
物体检测综述(第2112页)



提取全局语义信息的场景图
生成算法(第2214页)



融合弱监督目标定位的细粒
度小样本学习(第2226页)

学者观点

网络监督数据下的细粒度图像识别综述

魏秀参, 许玉燕, 杨健 2057

综述

海事监控视频舰船目标检测研究现状与展望

叶晨, 谔天洋, 肖灏灏, 陆海, 杨群慧 2078

深度学习行人检测方法综述

罗艳, 张重阳, 田永鸿, 郭捷, 孙军 2094

深度学习背景下视觉显著性物体检测综述

王自全, 张永生, 于英, 闵杰, 田浩 2112

图像分析和识别

图像增强对显著性目标检测的影响研究

郭继昌, 岳惠惠, 张怡, 刘迪, 刘晓雯, 郑司达 2129

混合高斯变分自编码器的聚类网络

陈华华, 陈哲, 郭春生, 应娜, 叶学义 2148

基于半监督对抗学习的图像语义分割

李志欣, 张佳, 吴璟莉, 马慧芳 2157

面向大姿态人脸识别的正面化形变场学习

胡蓝青, 阚美娜, 山世光, 陈熙霖 2171

融合时空域特征的人脸表情识别

陈拓, 邢帅, 杨文武, 金剑秋 2185

非局部注意力双分支网络的跨模态赤足足迹检索

鲍文霞, 茅丽丽, 王年, 唐俊, 杨先军, 张艳 2199

提取全局语义信息的场景图生成算法

段静雯, 闵卫东, 杨子元, 张煜, 陈鑫浩, 杨升宝 2214

融合弱监督目标定位的细粒度小样本学习

贺小箭, 林金福 2226

结合双注意力机制的道路裂缝检测

张志华, 温亚楠, 慕号伟, 杜小平 2240

融合神经网络的布料碰撞检测算法

靳雁霞, 马博, 贾瑶, 陈治旭, 芦烨 2251

双分支特征融合网络的步态识别算法

徐硕, 郑锋, 唐俊, 鲍文霞 2263

图像理解和计算机视觉

问题引导的空间关系图推理视觉问答模型

兰红, 张蒲芬 2274

结合扰动约束的低感知性对抗样本生成方法

王杨, 曹铁勇, 杨吉斌, 郑云飞, 方正, 邓小桐 2287

计算机图形学

动态书法墨迹的可回溯感评测

律睿悠, 梅莉琳, 昃跃峰, 晏涛 2300

CONTENTS

JOURNAL OF IMAGE AND GRAPHICS



Review of deep learning based salient object detection(P2112)



Global semantic information extraction based scene graph generation algorithm(P2214)



Weakly-supervised object localization based fine-grained few-shot learning(P2226)

Scholar View

Review of webly-supervised fine-grained image recognition

Wei Xiushen, Xu Yuyan, Yang Jian 2057

Review

Maritime surveillance videos based ships detection algorithms: a survey

Ye Chen, Lu Tianyang, Xiao Yuhao, Lu Hai, Yang Qunhui 2078

An overview of deep learning based pedestrian detection algorithms

Luo Yan, Zhang Chongyang, Tian Yonghong, Guo Jie, Sun Jun 2094

Review of deep learning based salient object detection

Wang Ziquan, Zhang Yongsheng, Yu Ying, Min Jie, Tian Hao 2112

Image Analysis and Recognition

The analysis of image enhancement on salient object detection

Guo Jichang, Yue Huihui, Zhang Yi, Liu Di, Liu Xiaowen, Zheng Sida 2129

A Gaussian mixture variational autoencoder based clustering network

Chen Huahua, Chen Zhe, Guo Chunsheng, Ying Na, Ye Xueyi 2148

Semi-supervised adversarial learning based semantic image segmentation

Li Zhixin, Zhang Jia, Wu Jingli, Ma Huifang 2157

Large pose face recognition with morphing field learning

Hu Lanqing, Kan Meina, Shan Shiguang, Chen Xilin 2171

Spatio-temporal features based human facial expression recognition

Chen Tuo, Xing Shuai, Yang Wenwu, Jin Jianqiu 2185

Non-local attention dual-branch network based cross-modal barefoot footprint retrieval

Bao Wenxia, Mao Lili, Wang Nian, Tang Jun, Yang Xianjun, Zhang Yan 2199

Global semantic information extraction based scene graph generation algorithm

Duan Jingwen, Min Weidong, Yang Ziyuan, Zhang Yu, Chen Xinhao, Yang Shengbao 2214

Weakly-supervised object localization based fine-grained few-shot learning

He Xiaojian, Lin Jinfu 2226

Dual attention mechanism based pavement crack detection

Zhang Zhihua, Wen Yanan, Mu Haowei, Du Xiaoping 2240

Neural network fusion based cloth collision detection algorithm

Jin Yanxia, Ma Bo, Jia Yao, Chen Zhixu, Lu Ye 2251

Dual branch feature fusion network based gait recognition algorithm

Xu Shuo, Zheng Feng, Tang Jun, Bao Wenxia 2263

Image Understanding and Computer Vision

Question-guided spatial relation graph reasoning model for visual question answering

Lan Hong, Zhang Pufen 2274

A perturbation constraint related weak perceptual adversarial example generation method

Wang Yang, Cao Tieyong, Yang Jibin, Zheng Yunfei, Fang Zheng, Deng Xiaotong 2287

Computer Graphics

Traceability evaluation of flowing calligraphy strokes

Lyu Ruimin, Mei Lilin, Ze Yuefeng, Yan Tao 2300

中图法分类号: TP391 文献标识码: A 文章编号: 1006-8961(2022)07-2171-14

论文引用格式: Hu L Q, Kan M N, Shan S G and Chen X L. 2022. Large pose face recognition with morphing field learning. Journal of Image and Graphics, 27(07): 2171-2184(胡蓝青, 阚美娜, 山世光, 陈熙霖. 2022. 面向大姿态人脸识别的正面化形变场学习. 中国图象图形学报, 27(07): 2171-2184)[DOI:10.11834/jig.210011]

面向大姿态人脸识别的正面化形变场学习

胡蓝青, 阚美娜, 山世光*, 陈熙霖

1. 中国科学院计算技术研究所, 北京 100190; 2. 中国科学院大学计算机科学与技术学院, 北京 100049

摘要: **目的** 人脸识别已经得到了广泛应用, 但大姿态人脸识别问题仍未完美解决。已有方法或提取姿态鲁棒特征, 或进行人脸姿态的正面化。其中主流的人脸正面化方法包括 2D 回归生成和 3D 模型形变建模, 前者能够生成相对自然真实的人脸, 但会引入额外的噪声导致图像信息的扭曲; 后者能够保持原始的人脸结构信息, 但生成过程是基于物理模型的, 不够自然灵活。为此, 结合 2D 和 3D 方法的优点, 本文提出了基于由粗到细形变场的人脸正面化方法。**方法** 该形变场由深度网络以 2D 回归方式学得, 反映的是不同视角人脸图像像素之间的语义级对应关系, 可以类 3D 的方式实现非正面人脸图像的正面化, 因此该方法兼具了 2D 正面化方法的灵活性与 3D 正面化方法的保真性, 且借鉴分步渐进的思路, 本文提出了由粗到细的形变场学习框架, 以获得更加准确鲁棒的形变场。**结果** 本文采用大姿态人脸识别实验来验证本文方法的有效性, 在 MultiPIE(multi pose, illumination, expressions)、LFW(labeled faces in the wild)、CFP(celebrities in frontal-profile in the wild)、IJB-A(intelligence advanced research projects activity Janus benchmark-A)等 4 个数据集上均取得了比已有方法更高的人脸识别精度。**结论** 本文提出的基于由粗到细的形变场学习的人脸正面化方法, 综合了 2D 和 3D 人脸正面化方法的优点, 使人脸正面化结果的学习更加灵活、准确, 保持了更多有利于识别的身份信息。

关键词: 大姿态人脸识别; 人脸正面化; 可学习形变场; 由粗到细学习; 全卷积网络

Large pose face recognition with morphing field learning

Hu Lanqing, Kan Meina, Shan Shiguang*, Chen Xilin

1. Institute of Computing Technology, Chinese Academy of Sciences, Beijing 100190, China;

2. School of Computer Science and Technology, University of Chinese Academy of Sciences, Beijing 100049, China

Abstract: **Objective** Face recognition is currently challenging in the context of large variations in pose, expression, aging, lighting and occlusion. Pose variations tend to large non-planar face transformation among these factors. To address the pose variations, previous methods mainly attempt to extract pose invariant feature or frontalize non-frontal faces. Among them, the frontalization methods can release discriminative feature learning via pose variations elimination. There are mainly two kinds of face frontalization methods: 2D and 3D frontalization methods. 2D methods can generate more natural frontal faces but it may lose facial structural information, which is the key factor of identity discrimination. 3D methods can well preserve facial structural information, but are not so flexible. In summary, both 3D methods and 2D methods have information loss in the frontalized faces especially for large pose variations like invisible pixels in 3D morphable model or pixel aberrance in

收稿日期: 2021-01-18; 修回日期: 2021-04-22; 预印本日期: 2021-04-29

* 通信作者: 山世光 sgshan@ict.ac.cn

基金项目: 国家重点研发计划资助(A1802); 国家自然科学基金项目(61772496)

Supported by: National Key R&D Program of China (A1802); National Natural Science Foundation of China (61772496)

2D generative methods. **Method** We propose a novel coarse-to-fine morphing field network (CFMF-Net), combining both 2D and 3D face transformation methods to frontalize a non-frontal face image via the coarse-to-fine optimized morphing field for shifting each pixel. Thanks to the flexibility of 2D learning based methods and structure preservation of 3D morphable model-based methods, our proposed morphing learning method makes the learning process easier and reduces the probability of over-fitting. First, a coarse morphing field is learned to capture the major structure variation of single face image. Then, a residual module based facial information extraction is designed to promote the coarse morphing field of those output concatenated with the coarse morphing field to generate the final fine morphing field for face image input. The overall framework is for the pixel correspondences regression but not pixel values. The work ensures that all pixels in the frontalized face image are taken from the input non-frontal image, thus reducing information distortion to a large extent. Therefore, the identity information related to the input non-frontal face images are well preserved with favorable visual results, thus further facilitating the subsequent face recognition task. To achieve more accurate morphing field output, our design of the coarse-to-fine morphing field learning assures the robustness of learned morphing field and the residual complementing branch. **Result** To verify the effectiveness of our proposed work, extensive experiments on multi pose, illumination, expressions (MultiPIE), labeled faces in the wild (LFW), celebrities in frontal-profile in the wild (CFP) and intelligence advanced research projects activity Janus benchmark-A (IJB-A) datasets are carried out and the results are compared with other face transformation methods. Among these testing sets, MultiPIE, CFP and IJB-A datasets are all with full pose variation. In addition, IJB-A contains full pose variations as well as other complicated variations like low resolution and occlusion. The experiments follow the same training and testing protocol with previous works, i. e., training with both original and frontalized face images. For fair comparison, the commonly used LightCNN-29 is developed as the recognition model. Our method outperforms related works on the large pose testing protocol of MultiPIE and CFP and comparable performance on LFW and IJB-A. Additionally, our visualization results also show that our method can well preserve the identity information. Furthermore, the ablation study presents the feasibility of the coarse-to-fine framework in our CFMF-Net. In a word, the recognition accuracies and visualization results demonstrate that the proposed CFMF-Net can generate frontalized faces with identity information preserved and achieve higher large pose face recognition accuracy as well. **Conclusion** A coarse-to-fine morphing field learning framework frontalizes face images by shifting pixels to ensure the flexible learnability and identity information preservation. To improve its accuracy, the flexible learnability yields the network to optimize face frontalization objective without predefined 3D transformation rules. Moreover, the learned morphing field for each pixel makes the output frontal face shifted from the input image only, reducing the information loss. Simultaneously, the design of coarse-to-fine and residual architecture ensures more robust and accurate results further.

Key words: large pose face recognition; face frontalization; morphing field learning; coarse-to-fine learning; fully convolutional network

0 引言

人脸识别技术在各领域的广泛应用,为人们的生活带来了巨大便利。随着技术的发展,人脸识别的性能得到了极大提升,非极端姿态的人脸识别已经取得了良好效果,但是大姿态下的人脸识别仍然面临很大挑战。这是由于人脸在大姿态下会发生很强的非平面内形变,影响对人脸身份的判别。主流的针对大姿态人脸识别问题的方法分为两大类:第1类方法直接在原图上提取姿态鲁棒特征,第2类方法先将人脸进行正面化之后再提取特征。第1类方法

用于极端姿态人脸识别时,可以提取的特征非常有限,人脸识别性能明显降低。第2类方法先将人脸正面化再进行公共特征提取,即人脸正面化方法,可以提取出更多有效的判别特征。正面化方法分为2D生成方法和利用3D模型变换方法。2D生成方法通过一个网络直接回归出正面人脸图像,3D方法则是将人脸图像建模为3D模型,通过该模型算出原图与正面人脸的像素坐标对应关系,从而实现正面化。2D生成方法比基于3D模型的方法更加灵活,生成的人脸也更加自然。然而3D方法得到的正面化人脸图像能够保留更多的人脸身份信息。

结合2D生成和3D模型变换两种正面化方法的优点,本文提出了一种基于由粗到细形变场学习的人脸正面化方法(coarse-to-fine morphing field network, CFMF-Net),如图1所示。CFMF-Net通过学习形变场将任意人脸图像 I 正面化为图像 I^{est} 。该网络首先通过 F^s 提取人脸关键点特征 S ,并将 S 输入 F^c 以得到粗粒度形变场 C 。之后将 C 和 F^g 学到的细节特征 G 拼接在一起,输入 F^d 以得到细粒度形变场 D 。形变模块 T 将形变场 D 作用于输入图像 I 得到 I^{est} 。CFMF-Net通过拉近 I^{est} 与真实的正脸图像 I^{gt} 的距离来进行优化。其中形变场的值由下方的热力图表示,红色表示该像素点上的形变场向左,蓝色表示该像素点上的形变场向右,颜色明度越高则移动距离越大。此处形变场指正面人脸与输入人脸的像素点的位置对应关系,即非正面人脸图像的像素可以根据形变场进行重组得到对应的正面人脸图像。CFMF-Net通过一个深度网络以由粗到细的优化策略学习形变场,对输入人脸进行正面化。

本文采用的以形变场进行正面化的方式与利用3D人脸模型进行人脸正面化的方法类似,都能够通过像素点的移动来变换图像,保证正面化人脸图像中的像素点全部来源于原始图像。并且本文方法与2D回归方法类似,都是通过网络自动学习,而不是人为设计的规则。因此该方法兼具了2D正面化方法的灵活性与3D正面化方法的保真性。

目标形变场来自高维空间,这给网络的优化带来了不小的难度。因此本文借鉴分步渐进的优化思路,提出了由粗到细的形变场学习框架,以获得更加准确鲁棒的形变场。然而在学习粗粒度形变信息时,模型只留意了人脸的主要形变,会导致细节信息的丢失,因而增加了一个细节补充分支网络,以进一步保证预测出的形变场的准确性。

本文的主要贡献在于:1)采用2D回归的方式以类3D的行为对人脸进行正面化,结合了2D正面化方法的灵活性与3D正面化方法的保真性;2)由粗到细的学习方式提升了模型的易学习性。

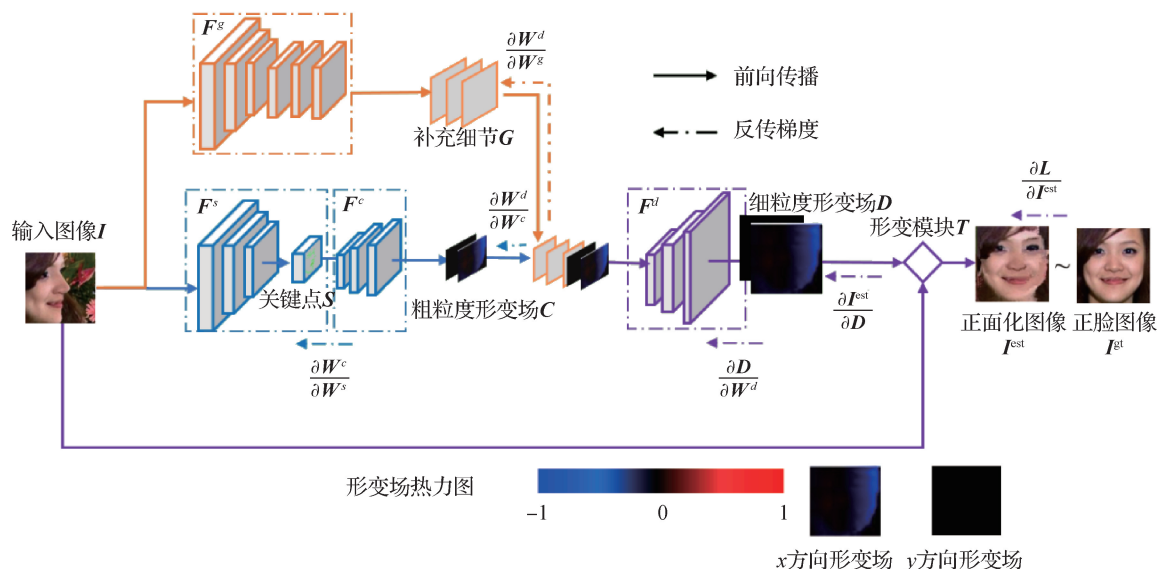


图1 基于由粗到细形变场学习的方法 CFMF-Net 流程图

Fig. 1 Overview of our CFMF-Net

1 相关工作

大姿态人脸识别方法分为直接在原图上提取姿态鲁棒特征和先将人脸正面化再提取特征两类。

直接提取姿态鲁棒特征的方法主要是将不同姿态的人脸图像都映射到一个公共的特征空间中。典

型相关分析(Li等,2009)是直接提取姿态鲁棒特征早期的经典方法,通过最大化两组不同姿态的图像的特征相关性,将不同姿态的特征映射到统一的空间中。然而该方法只保证了提取到的是不同姿态图像的公共特征,忽略了特征的判别能力。Sharma和Jacobs(2011)改进了典型相关分析,通过偏最小二乘法最小化同一个人所有姿态的图像的特征距离,

得到的特征不仅是姿态鲁棒的,且具有较好的判别能力。Zhang 等人(2013)给训练集中同一人所有姿态的图像设定同一张随机的人脸作为映射目标,以得到姿态鲁棒的具有良好判别能力的特征。多视角判别网络(Kan 等,2016)针对不同姿态的图像采用不同的特征映射,将多姿态的图像映射到公共特征空间中,得到了更准确的公共判别特征。深度网络的提出与发展进一步赋予了模型更强大的特征学习能力。基于深度学习的特征解耦方法(Peng 等,2017)首先利用深度网络提取出更准确的人脸表示,之后通过特征解耦与交叉重组得到姿态鲁棒特征。

这些直接提取姿态鲁棒特征的方法对非极端姿态的人脸识别已经有了不错的效果,但对极端姿态的人脸识别却效果有限。因为对这些姿态差异巨大的人脸图像直接提取公共特征会丢失很多对识别有用的信息。因此研究者提出了先将人脸转正,再进行人脸识别的人脸正面化方法,这些方法又分为2D正面化方法和3D模型正面化方法。图2展示了几种经典方法在MultiPIE数据集上的正面化结果。

2D人脸正面化方法直接通过一个编码器网络将不同姿态的人脸图像映射为正面姿态的图像。经典的方法(Zhu 等,2013; Kan 等,2014)是用渐进式学习的方式对侧面人脸进行逐步的姿态调整,以映射到正面人脸。随着生成对抗网络(generative adversarial network, GAN)(Goodfellow 等,2014)的提出,很多方法借助GAN强大的分布拟合能力生成各种姿态的人脸,包括正脸。相比于通过回归生成人脸的方法,基于GAN的方法生成的人脸图像更加逼真。在Luan 等人(2017)方法中,由特征提取器得到的身份特征和指定的姿态信息一起输入GAN中,

以生成多姿态的人脸图像。Yin 等人(2017)提出了另一个更精细的基于GAN的方法,给予GAN更多的信息,即3D可变形模型的系数,得到保留了更多原始信息的正面人脸图像。Huang 等人(2017)同时兼顾整张人脸和人脸局部图像块的逼真程度,使生成的人脸图像保留了更多的细节。Zhang 等人(2019)认为更大姿态的人脸更难以识别与正面化,因此在通过GAN正面化人脸的训练过程中对难样本采用更大的训练权重。Rong 等人(2020)通过特征级和图像级两种GAN判别器,加强GAN正面化人脸的效果。Luan 等人(2020)在GAN判别器中加入自注意力机制保持人脸图像的几何结构,令人脸正面化更加真实。

3D人脸正面化方法通过建立人脸图像的3D模型将人脸映射到正面姿态。相比于2D方法,3D人脸正面化方法能保留更多的人脸结构信息。早期的经典方法,3D通用弹性模型(Prabhu 等,2011)和基于视角的主动外观模型(Asthana 等,2011)等直接利用3D模型进行人脸姿态变换。这些方法通过将2D图像映射到3D坐标上,再投影到任意的角度,以生成相应姿态的人脸。更直接的方法是计算侧面人脸图像到其正面人脸图像的像素点的位置对应关系,即形变场,再用该形变场进行图像变换。Li 等人(2012)用从训练集得到的正面化形变场的线性组合来正面化测试集人脸图像。而这些3D方法都不能处理姿态变化引起的自遮挡,如图2(b)所示。Ding 等人(2015)在3D模型变换的基础上,利用人脸的对称性填补遮挡部分,但生成的人脸依然存在严重的失真,如图2(c)所示。Hu 等人(2017)提出了一种利用全连接网络自动回归正面化形变场的方法,生成了更逼真并保留更多原始信息的正面

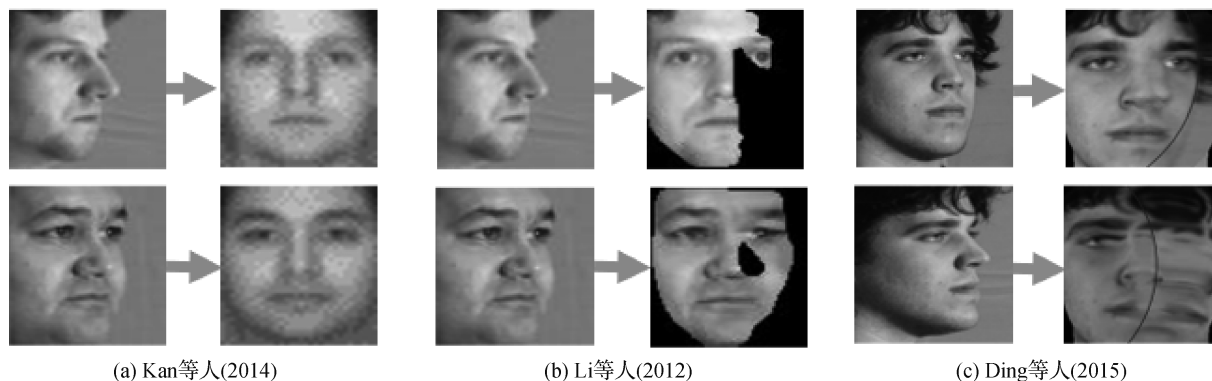


图2 3种经典方法在MultiPIE数据集上的正面化结果

Fig. 2 Visualization results of three methods ((a) Kan et al. (2014); (b) Li et al. (2012); (c) Ding et al. (2015))

人脸。Cao 等人(2018)提出一种结合了 3D 模型和 GAN 的方法,首先通过形变场得到一个初始的正脸图像,再通过 GAN 进行图像调整,最终得到足够逼真且身份保持的正面人脸。

综上所述,人脸正面化方法相比于直接提取姿态鲁棒特征的方法能够提取出更有效的公共判别特征。正面化方法中,2D 方法比 3D 方法更加灵活,生成的人脸也更加自然。3D 方法得到的正面化人脸图像能够保留更多的人脸身份信息。

2 本文方法

如图 1 所示,本文提出的 CFMF-Net 主要由可学习形变场网络 F 和用形变场进行正面化的模块 T 两部分组成。网络 F 的输入为原始人脸图像 I ,其输出为正面化 I 的形变场 D 。 T 的输入为原始图像 I 和形变场 D ,其输出为正面化后的图像 I^{est} 。

可学习形变场网络 F 通过渐进式的方式学习形变场,即先学习粗粒度形变场以捕捉人脸结构的主要形变,在此基础上再学习细粒度形变场来精修细节上的形变。因此,网络 F 主要包含粗粒度形变场网络 F^c 和细粒度形变场网络 F^d 两部分。具体来讲, F^c 首先学习人脸关键点,再解码出粗粒度形变场。 F^d 进一步完善粗粒度形变场,得到与原图同分辨率的细粒度形变场,其输入包含 F^c 的输出与一个分支网络 F^s 从原图学到的补充细节两部分。

CFMF-Net 通过学习到的形变场对图像进行变换,因而其输出图像的像素值都是来自于原图,保留了更多的身份信息,减少了额外噪声的引入。相比于 2D 方法通过回归像素值生成正脸图像,本文方法通过学习形变场进行正面化,从而限制了正面化图像中的像素均来自于原图,更好地保持了原始信息。相比于 3D 方法基于 3D 模型规则计算形变场,本文方法得到的形变场是基于学习得到的,从而能够得到更逼真的正面化结果。

2.1 形式化

本文方法利用成对的人脸数据 $\{(I_1, I_1^{\text{gt}}), \dots, (I_n, I_n^{\text{gt}})\}$ 进行训练,其中 $I_k \in \mathbf{R}^{h \times w}$ 与 $I_k^{\text{gt}} \in \mathbf{R}^{h \times w}$ 分别代表输入人脸图像和其对应的正面人脸图像。CFMF-Net 通过一个深度网络学习形变场来正面化输入图像 I_k ,得到估计的正面人脸图

像 I_k^{est} ,其目标为最小化真实正脸图像和估计得到的正脸图像的差别。即

$$\min_W \frac{1}{n} \sum_{k=1}^n \|I_k^{\text{est}} - I_k^{\text{gt}}\|_2^2 \quad (1)$$

式中, W 为整个模型的可学习参数。

更具体地讲, I_k^{est} 是由形变模块 T 将形变场网络 F 得到的形变场 $F(I_k)$ 作用于输入图像 I_k 形变而来。即

$$I_k^{\text{est}} = T(I_k, F(I_k)) \quad (2)$$

2.1.1 形变场学习网络

根据形变场的定义,可以得到输出图像 I_k^{est} 上位于 (i, j) 位置的像素点 $I_{k,i,j}^{\text{est}}$ 的形变场为

$$D_{k,i,j} = (\Delta_{k,i,j}^h, \Delta_{k,i,j}^w) \quad (3)$$

其表示该像素点取自输入图像 I_k 的 $(i + \Delta_{k,i,j}^h, j + \Delta_{k,i,j}^w)$ 位置。如果 $(i + \Delta_{k,i,j}^h, j + \Delta_{k,i,j}^w)$ 超出图像空间范围,则分别对 h 和 w 取余。为了表示简单,后文仍将最终坐标写做 $i + \Delta_{k,i,j}^h, j + \Delta_{k,i,j}^w$ 。

粗粒度形变场网络 F^c 和细粒度形变场 F^d 是 CFMF-Net 的两个重要组成部分。

F^c 的目标为学习输入到输出图像的人脸结构的主要变化,其目标为得到一个大小为 $\frac{h}{4} \times \frac{w}{4}$ 的粗粒度形变场 C_k ,即

$$S_k = F^s(I_k) \quad (4)$$

$$C_k = F^c(S_k) \quad (5)$$

式中, F^s 和 F^c 是两个连接在一起的卷积网络,其参数分别为 W^s 和 W^c 。 $S_k \in \mathbf{R}^{68 \times 2}$ 为 68 个稀疏人脸关键点的位置表示,作为人脸结构鲁棒特征表示用来指导粗粒度形变场的学习,而学得形变场 C_k 将作为学习大小为 $h \times w$ 的细粒度形变场的中间表示,为细粒度形变场学习打下良好基础。

C_k 建模了输入到输出人脸图像的主要形变,但 C_k 忽略了细节的变化,因此还需要进一步细化。在 CFMF-Net 中,分支网络 F^s 用来提取原始图像 I_k 的细节特征 $G_k = F^s(I_k)$,其中 F^s 的参数为 W^s 。之后,将 C_k 与 G_k 拼接在一起,输入到细粒度形变场网络 F^d 中,得到与原图分辨率大小相同的细粒度形变场 $D_k \in \mathbf{R}^{h \times w}$ 。即

$$D_k = F^d([C_k, G_k]) \quad (6)$$

式中, F^d 为反卷积网络,可以对粗粒度形变场进行上采样,其参数为 W^d 。

2.1.2 形变模块

得到形变场 D_k 之后,形变模块 T 将 D_k 作用于原图 I_k ,得到正面化后的图像 I_k^{est} 。即

$$I_k^{\text{est}} = T(I_k, D_k) \quad (7)$$

T 通过形变场 D_k 将原始图像 I_k 的像素点进行重组得到 I_k^{est} ,这个过程没有可学习参数。如果 $D_{k,i,j} = (\Delta_{k,i,j}^h, \Delta_{k,i,j}^w)$ 是整数, I_k^{est} 中位于坐标 (i, j) 的像素就直接取自于原图 I_k 中位于坐标 $(i + \Delta_{k,i,j}^h, j + \Delta_{k,i,j}^w)$ 上的像素。即

$$I_{k,i,j}^{\text{est}} = I_{k,i+\Delta_{k,i,j}^h, j+\Delta_{k,i,j}^w} \quad (8)$$

一般情况下, $D_{k,i,j}$ 是实数(为了方便求导,本文中并不会限制它是整数)。此时, $I_{k,i,j}^{\text{est}}$ 像素值为 $I_k(i + \Delta_{k,i,j}^h, j + \Delta_{k,i,j}^w)$ 邻近 4 个像素点像素值的双线性插值。令 $\tilde{i} = i + \Delta_{k,i,j}^h, \tilde{j} = j + \Delta_{k,i,j}^w$, 则

$$I_{k,i,j}^{\text{est}} = (1 - \lfloor \tilde{i} \rfloor - \tilde{i}) \times (1 - \lfloor \tilde{j} \rfloor - \tilde{j}) \times I_{k,\lfloor \tilde{i} \rfloor, \lfloor \tilde{j} \rfloor} + (1 - \lfloor \tilde{i} \rfloor - \tilde{i}) \times (1 - \lceil \tilde{j} \rceil - \tilde{j}) \times I_{k,\lfloor \tilde{i} \rfloor, \lceil \tilde{j} \rceil} + (1 - \lceil \tilde{i} \rceil - \tilde{i}) \times (1 - \lfloor \tilde{j} \rfloor - \tilde{j}) \times I_{k,\lceil \tilde{i} \rceil, \lfloor \tilde{j} \rfloor} + (1 - \lceil \tilde{i} \rceil - \tilde{i}) \times (1 - \lceil \tilde{j} \rceil - \tilde{j}) \times I_{k,\lceil \tilde{i} \rceil, \lceil \tilde{j} \rceil} \quad (9)$$

容易看出,式(8)是式(9)的特殊情况。从式(8)和式(9)可以看出,正面化图像 I_k^{est} 中所有的像素点都来自原图 I_k 某一个像素点或者由 4 个邻近点加权得到。因此, I_k^{est} 极大保留了 I_k 中的原始信息。

2.1.3 整体训练目标

本文方法通过端到端的学习方式优化整个形变场网络 CFMF-Net 为 $\{F^s, F^c, F^d, F^g\}$, 目标为正面化后的人脸图像 I_k^{est} 与真实的正脸图像 I_k^{gt} 尽量相同。即

$$L = \min_W \frac{1}{n} \sum_{k=1}^n \|I_k^{\text{est}} - I_k^{\text{gt}}\|_2^2 = \min_{W^s, W^c, W^d, W^g} \frac{1}{n} \times \sum_{k=1}^n \|T(I_k, F^d([F^c(F^s(I_k)), F^g(I_k)])) - I_k^{\text{gt}}\|_2^2 \quad (10)$$

2.2 优化过程

为了加快 CFMF-Net 的收敛,首先预训练 CFMF-Net 每个模块,得到一个较好的初始化参数,再以式(10)为目标进行端到端的训练。

2.2.1 预训练

如前所述,粗粒度形变场学习中的 F^s 用来学习人脸关键点位置 S_k 。此处用人脸关键点对 F^s 进行优化(若无标定的关键点可省略该步)。即

$$L^s = \min_{W^s} \frac{1}{n} \sum_{k=1}^n \|F^s(I_k) - S_k^{\text{gt}}\|_2^2 \quad (11)$$

式中, S_k^{gt} 是人工标注的人脸关键点,如图 3 所示。

F^s 通过梯度下降进行优化,梯度为 $\frac{\partial L^s}{\partial W^s}$ 。



图3 人脸关键点示例

Fig. 3 Exemplars of facial landmarks

粗粒度形变场网络 F^c 以人脸关键点位置 S_k 为输入,学习粗粒度形变场 C_k 。本文方法借助事先计算得到的粗粒度形变场 $\hat{C}_k^{\text{gt}}$ 对 F^c 进行初始化。即

$$L^c = \min_{W^c} \frac{1}{n} \sum_{k=1}^n \|C_k - \hat{C}_k^{\text{gt}}\|_2^2 = \min_{W^c} \frac{1}{n} \sum_{k=1}^n \|F^c(S_k) - \hat{C}_k^{\text{gt}}\|_2^2 \quad (12)$$

式中, $\hat{C}_k^{\text{gt}}$ 的大小为 $\frac{h}{4} \times \frac{w}{4}$ 。为了加速预处理过程,

对于同一姿态的所有人脸图像,只取出一幅代表性的人脸图像,计算得到一个统一的 $\hat{C}_k^{\text{gt}}$ 。给定人脸图像 I_k 和其对应正面人脸 I_k^{est} ,借助这对图像的关键点,利用薄板样条插值(thin plate spline, TPS)(Bookstein, 1989)粗略地估算出一个正面化形变场。但 TPS 无法填补自遮挡部分,因此,为了解决自遮挡问题,本文方法进一步利用人脸对称部分补齐遮挡的像素点,从而得到最终估算的 $\hat{C}_k^{\text{gt}}$,具体实现过程如图 4 所示。同样, F^c 也通过梯度下降进行优化,梯度为 $\frac{\partial L^c}{\partial W^c}$ 。

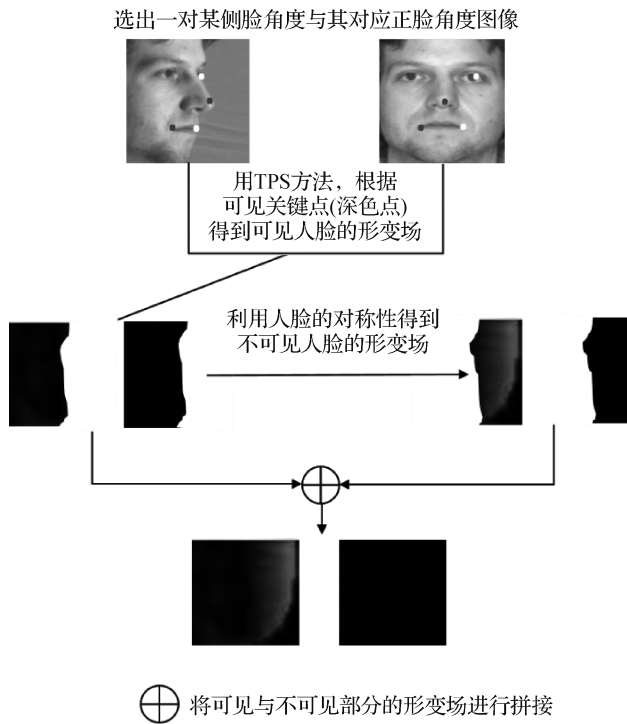
类似地,借助事先计算得到的细粒度正面化形变场 $\hat{C}_k^{\text{gt}}$ 可以对细粒度形变场网络 F^d 和细节分支网络 F^g 一同进行初始化训练。即

$$L^d = \min_{W^g, W^d} \frac{1}{n} \sum_{k=1}^n \|D_k - \hat{D}_k^{\text{gt}}\|_2^2 =$$

$$\min_{W^g, W^d} \frac{1}{n} \sum_{k=1}^n \|F^d([C_k, F^g(I_k)]) - \hat{D}_k^{\text{gt}}\|_2^2 \quad (13)$$

与 $\hat{C}_k^{\text{gt}}$ 类似, $\hat{D}_k^{\text{gt}}$ 也同样由 TPS 得到,只是其大小为 $h \times w$,可以看做是 $\hat{C}_k^{\text{gt}}$ 的上采样。同样, F^d 和 F^g 也通过梯度下降进行优化,梯度分别为 $\frac{\partial L^d}{\partial W^d}$ 和

$$\frac{\partial L^d}{\partial W^g}。$$

图 4 CFMF-Net 预训练时粗略估计 $\hat{C}_k^{est}$ 过程Fig. 4 The estimating process of $\hat{C}_k^{est}$ during pretraining

2.2.2 端到端调优

在预训练的基础上, CFMF-Net 以式 (10) 为目标对网络进行端到端的优化。

由式 (10), 优化目标 L 关于 I_k^{est} 的导数为

$$\frac{\partial L}{\partial I_{k,i,j}^{est}} = 2(I_{k,i,j}^{est} - I_{k,i,j}^{gt}) \quad (14)$$

由式 (9), I_k^{est} 关于形变场 D_k 的导数为对每个像素点进行求导, 分为长方向和宽方向上的形变场两部分 $\left(\frac{\partial I_{k,i,j}^{est}}{\partial \Delta_{k,i,j}^h}, \frac{\partial I_{k,i,j}^{est}}{\partial \Delta_{k,i,j}^w}\right)$, 具体为

$$\begin{aligned} \frac{\partial I_{k,i,j}^{est}}{\partial \Delta_{k,i,j}^h} &= \frac{\partial I_{k,i,j}^{est}}{\partial \tilde{i}} \times \frac{\partial \tilde{i}}{\partial \Delta_{k,i,j}^h} = \frac{\partial I_{k,i,j}^{est}}{\partial \tilde{i}} = \\ &= (|\lfloor \tilde{j} \rfloor - \tilde{j}| - 1) \times (I_{k,\lfloor \tilde{i} \rfloor, \lfloor \tilde{j} \rfloor} - I_{k,\lceil \tilde{i} \rceil, \lfloor \tilde{j} \rfloor}) + \\ &+ (|\lceil \tilde{j} \rceil - \tilde{j}| - 1) \times (I_{k,\lfloor \tilde{i} \rfloor, \lceil \tilde{j} \rceil} - I_{k,\lceil \tilde{i} \rceil, \lceil \tilde{j} \rceil}) \quad (15) \end{aligned}$$

$$\begin{aligned} \frac{\partial I_{k,i,j}^{est}}{\partial \Delta_{k,i,j}^w} &= \frac{\partial I_{k,i,j}^{est}}{\partial \tilde{j}} \times \frac{\partial \tilde{j}}{\partial \Delta_{k,i,j}^w} = \frac{\partial I_{k,i,j}^{est}}{\partial \tilde{j}} = \\ &= (|\lfloor \tilde{i} \rfloor - \tilde{i}| - 1) \times (I_{k,\lfloor \tilde{i} \rfloor, \lfloor \tilde{j} \rfloor} - I_{k,\lceil \tilde{i} \rceil, \lfloor \tilde{j} \rfloor}) + \\ &+ (|\lceil \tilde{i} \rceil - \tilde{i}| - 1) \times (I_{k,\lfloor \tilde{i} \rfloor, \lceil \tilde{j} \rceil} - I_{k,\lceil \tilde{i} \rceil, \lceil \tilde{j} \rceil}) \quad (16) \end{aligned}$$

整个 CFMF-Net 网络参数 $\{W^d, W^g, W^c, W^s\}$ 通过梯度下降进行优化, 对应每个模块的梯度为

$$\frac{\partial L}{\partial W^d} = \frac{\partial L}{\partial I_k^{est}} \times \frac{\partial I_k^{est}}{\partial D_k} \times \frac{\partial D_k}{\partial W^d}$$

$$\begin{aligned} \frac{\partial L}{\partial W^g} &= \frac{\partial L}{\partial W^d} \times \frac{\partial W^d}{\partial W^g} \\ \frac{\partial L}{\partial W^c} &= \frac{\partial L}{\partial W^d} \times \frac{\partial W^d}{\partial W^c} \\ \frac{\partial L}{\partial W^s} &= \frac{\partial L}{\partial W^d} \times \frac{\partial W^d}{\partial W^c} \times \frac{\partial W^c}{\partial W^s} \quad (17) \end{aligned}$$

3 实验

为验证本文方法对大姿态人脸识别问题的有效性, 在 4 个代表性大姿态人脸识别数据集上进行实验, 包括通用人脸识别数据集 LFW (labeled faces in the wild)、包含更多更极端姿态变化的数据集 MultiPIE (multi pose, illumination, expressions)、CFP (celebrities in frontal-profile in the wild) 和 IJB-A (intelligence advanced research projects activity janus benchmark-A)。

3.1 数据集与实验设置

在 MultiPIE 数据集 (Sim 等, 2003) 上进行可控场景下的姿态人脸识别实验, 在 300 W-LP (Zhu 等, 2015)、Webface (Yi 等, 2014)、LFW (Huang 和 Learned-Miller, 2014)、CFP (Sengupta 等, 2016) 和 IJB-A (Klare 等, 2015) 上进行非可控场景下的姿态人脸识别实验。在所有实验中, 首先通过 CFMF-Net 进行人脸正面化, 之后通过一个人脸识别网络进行人脸识别。其中, 300 W-LP 为 CFMF-Net 网络的训练集, Webface 为人脸识别训练集, LFW、CFP 和 IJB-A 为人脸识别测试集。训练集和测试集的设置情况如表 1 所示。实验时, 通过裁剪缩放, 所有的人脸图像调整至 128×128 像素, 像素值归一化至 $[-1, 1]$, 图像坐标值归一化到 $[0, 1]$, 形变场归一到 $[-1, 1]$ 。图 5 展示了不同实验的 CFMF-Net 网络结果。接下来具体介绍实验中的数据集。

MultiPIE 数据集 (Sim 等, 2003) 是最常用的可控场景下的姿态人脸识别数据集, 包含 337 个人在不同姿态、光照和表情下的照片。实验采用与大姿态人脸识别的代表性工作 (Cao 等, 2018) 相同的实验设置, 即取前 200 个人的所有图像进行人脸正面化和识别的训练, 剩下 137 个人的所有图像进行测试。在测试阶段, 采用这 137 个人的正面姿态、光照和中性表情的照片作为注册集 (gallery), 剩下 72 000 张照片作为查询集 (probe)。与大多数对比

表 1 训练集和测试集的设置说明
Table 1 Overview of training and testing datasets

测试集	条件是否可控	是否包含大姿态	训练集	
			正面化	识别
MultiPIE	是	是	MultiPIE	MultiPIE
LFW	否	否	300 W-LP	Webface
CFP	否	是	300 W-LP	Webface
IJB-A	否	是	300 W-LP	Webface + IJB-A

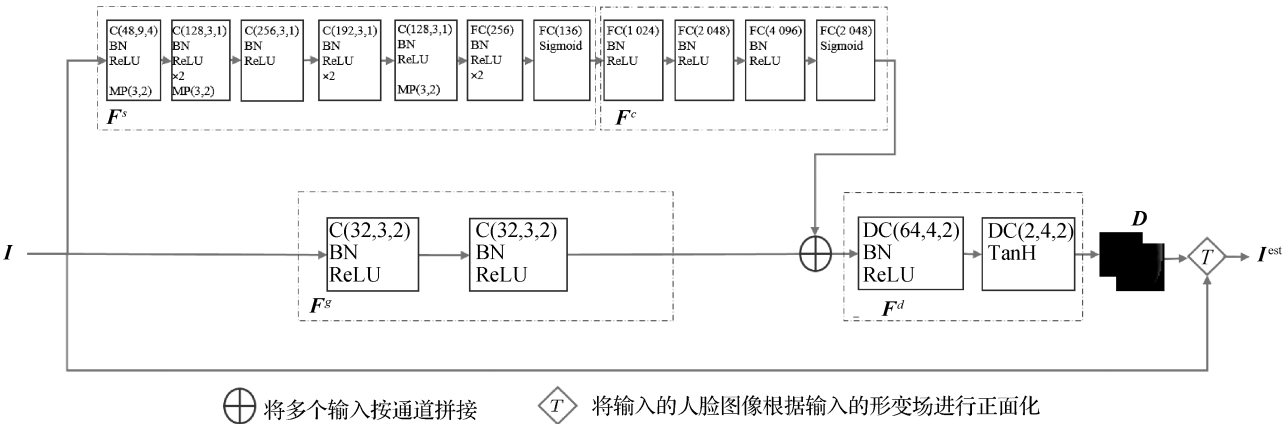


图 5 CFMF-Net 网络结构

Fig. 5 Architecture of CFMF-Net

方法相同,在 MultiPIE 的实验中,本方法采用 Light-CNN-29(Wu 等,2018)作为识别网络。

LFW(Huang 和 Learned-Miller, 2014) 和 CFP(Sengupta 等,2016) 是两个经典的非可控场景下的人脸识别数据集,通常用来测试人脸识别方法的性能。LFW 包含 13 233 幅采集自网络的人脸图像,其中通常用于人脸识别测试的部分为 3 000 对来自于同一人的图像与 3 000 对来自于不同人的图像。CFP 包含来自 500 人的 7 000 幅图像,其中每个人都有 10 幅准正面(小于 10°)图像和 4 幅大姿态(大于 10°)的图像。本文实验中,LFW 和 CFP 用来进行人脸验证实验。在 LFW 上的测试指标为人脸验证准确率 ACC(accuracy) 与 ROC(receiver operating characteristic curve) 曲线下的面积 AUC(area under the curve)。在 CFP 上的测试包含正脸—正脸图像对(frontal-frontal, FF) 和正脸—侧脸图像对(frontal-profile, FP) 两部分,其测试指标为人脸验证准确率 ACC。同样,在 LFW 和 CFP 实验中,本文方法用 LightCNN-29(Wu 等,2018) 作为识别网络。

IJB-A(Klare 等,2015) 是更大的不可控场景下的人脸识别数据集,主要用来测试大姿态人脸识别方法的性能。IJB-A 包含很多极端姿态和光照条件下的人脸图像,相比于前面介绍的测试数据集,更具有挑战性。IJB-A 包含来自 500 人的 5 396 幅网络图像和 20 412 幅截取自网络视频的图像。其测试协议为 10 折交叉验证,每次划分出 333 人的图像作为训练集,剩余 167 人的图像作为测试集,最终的准确率为 10 次实验的平均准确率。在多数方法中,首先在一个更大数据集(如 Webface) 上训练一个识别模型,再用 IJB-A 的小训练集进行微调(Klare 等,2015)。相比于之前介绍的数据集,IJB-A 上的测试不再是单一图像的对比,而是图像集合之间的对比。测试包含人脸验证和人脸识别两部分。人脸验证的指标为在某个指定错误接受率(false accept rate, FAR) 下的正确接受率(true accept rate, TAR)。人脸识别通常为闭集测试,指标为第 1 名准确率和前 5 名准确率。在之前的方法中,IJB-A 上的测试没有统一的训练集和训练网络结构,为了与之前的方法公

平比较,本文方法采用了两个不同的人脸识别网络,分别为 Fast AlexNet 和 LightCNN-29(Wu 等,2018)。其中, Fast AlexNet 是对 AlexNet 进行优化后得到的模型,与大多数已有方法的模型能力相当,但收敛速度更快,具体结构如表 2 所示。

表 2 Fast AlexNet 网络结构
Table 2 Architecture of Fast AlexNet

模块	结构
1	Conv $48 \times 9 \times 9$, S4, BatchNorm/ReLU/maxpool 3×3 , S2
2	Conv $128 \times 3 \times 3$, S1, BatchNorm/ReLU
3	Conv $128 \times 3 \times 3$, S1, BatchNorm/ReLU/maxpool 3×3 , S2
4	Conv $256 \times 3 \times 3$, S1, BatchNorm/ReLU
5	Conv $192 \times 3 \times 3$, S1, BatchNorm/ReLU
6	Conv $192 \times 3 \times 3$, S1, BatchNorm/ReLU
7	Conv $128 \times 3 \times 3$, S1, BatchNorm/ReLU/maxpool 3×3 , S2
8	FC 4 096/BatchNorm/ReLU
9	FC 2 048/BatchNorm/ReLU

注:S 为卷积核的滑动步长 stride。

300 W-LP 是人脸姿态增广方法(Zhu 等,2015)对 300 W 数据集(Sagonas 等,2016)增广后构建的增广数据集,包含 122 450 幅图像。实验中,每个正面图像与对应的所有增广侧面图像作为人脸正面化训练集,而该准正面图像就作为训练目标(训练数据列表见 [https://github.com/whobefore/MF-Net/tree/master/Data/300 W-LP](https://github.com/whobefore/MF-Net/tree/master/Data/300%20W-LP))。并且,300 W-LP 中的每幅人脸图像包含 68 个人脸关键点标注,可以作为 S_k^{gt} 对 F^s 进行预训练。

Webface(Yi 等,2014)是一个通用人脸识别训练集,包含来自 10 575 个人的 494 414 幅图像。实验中,使用 Webface 训练非可控条件下的人脸识别模型。

3.2 实验结果

在 MultiPIE、LFW 和 CFP 数据集上,本文提出的 CFMF-Net 与多种方法进行实验对比,包括多任务学习方法(Yin 和 Liu,2018)以及与本文方法同为图像生成类的基于 GAN 的方法(Luan 等,2017;Yin 等,2017;Zhao 等,2018a,b;Cao 等,2018)。其中 Luan 等人(2017)的方法是一种直接基于 GAN 的 2D 人脸正面化方法。Yin 等人(2017)在 DR-GAN 的基础上进一步抽取了 3DMM(3D morphable model)的系数作为特征,更好地保持了人脸结构信息。Cao

等人(2018)首先将形变场作用于原图得到正面化人脸,再以此为中间结果做进一步调整。

在 IJB-A 数据集上,本文方法与不同类型的方法进行了对比,包括特征解耦方法(Crosswhite 等,2017;Yang 等,2017;Zhao 等,2017)、人脸增广方法(Zhu 等,2016;Masi 等,2017;Chang 等,2017)和人脸正面化方法(Luan 等,2017;Yin 等,2017;Zhao 等,2018a;Cao 等,2018)。

值得一提的是,2019 年以后出现的方法多为通用人脸识别方法,极少针对大姿态人脸识别这一特定问题专门研究,本文与 ArcFace(Deng 等,2019)采用 ResNetSE50 网络结构(网络能力与本文网络差不多)的版本(<https://github.com/TreBlE/Insight-Face-Pytorch>)进行比较。

在 MultiPIE 数据集上实验结果如表 3 所示。可以看出,本文方法得到了比对比方法更好的结果,尤其是在 $75^\circ \sim 90^\circ$ 的大姿态人脸下。同样是利用深度网络自动学习形变场,Hu 等人(2017)的方法由于结构简单,在表 3 所采用的更复杂的数据集上并不能收敛,且在较小分辨率 32×32 的 MultiPIE 上的平均性能比 CFMF-Net 低 0.4%。

表 3 不同方法在 MultiPIE 数据集上的识别率
Table 3 Face recognition accuracy of different methods on MultiPIE dataset

方法	人脸姿态角度						/%
	$\pm 90^\circ$	$\pm 75^\circ$	$\pm 60^\circ$	$\pm 45^\circ$	$\pm 30^\circ$	$\pm 15^\circ$	
Luan 等人(2017)	-	-	83.20	86.20	90.10	94.00	
Huang 等人(2017)	64.64	77.43	87.72	95.38	98.06	98.68	
Yin 等人(2018)	76.96	87.83	92.07	90.34	98.01	99.19	
Zhao 等人(2018b)	86.73	95.21	98.37	98.81	99.48	99.64	
Cao 等人(2018)	92.32	96.40	99.14	99.88	99.98	99.99	
本文	94.00	97.76	99.36	99.94	99.98	100.00	

注:加粗字体表示各列最优结果,“-”表示原方法未报告结果。

在 LFW 和 CFP 上的实验结果如表 4 和表 5 所示。可以看出,本文方法在正面人脸居多的测试中与当前最好方法的性能相当,包括采用更大训练集的 Deng 等人(2019)方法。从表 4 可以看到,在 LFW 数据集上,本文方法得到了保持原始信息的正面化人脸。从如表 5 可以看到,本文方法在正脸—

表 4 不同方法在 LFW 数据集上的人脸验证
准确率 ACC 和 AUC
Table 4 Face verification accuracy and area under curve of
different methods on LFW dataset

方法	ACC	AUC
Zhu 等人(2015)	96.25	99.39
Yin 等人(2017)	96.42	99.45
Cao 等人(2018)	99.41	99.92
本文	99.02	99.92

注:加粗字体表示各列最优结果。

表 5 不同方法在 CFP 数据集上的人脸验证准确率 ACC
Table 5 Face verification accuracy of
different methods on CFP dataset

方法	正脸—正脸	正脸—侧脸	平均
Luan 等人(2017)	97.84	93.41	95.63
Yin 等人(2018)	97.79	94.39	96.09
Zhao 等人(2018a)	99.44	93.10	96.27
Deng 等人(2019)	99.62	95.04	97.33
本文	99.34	95.17	97.26

注:加粗字体表示各列最优结果。

侧脸的识别上取得了更好性能,表明本文方法的正

面化对侧面人脸识别起到了重要作用。

在 IJB-A 数据集上的实验结果如表 6 所示。在人脸正面化类方法中,本文方法与当前最好的方法效果相当。表 6 中,本文方法 CFMF-Net1 是以最大化真实正面人脸与生成正面人脸的相似度为目标,学习原图与正面化图像的形变场,通过重组原图像素点得到正面化的图像,保证生成图像的所有像素都来自原图。Masi 等人(2017)、Luan 等人(2017)和 Yin 等人(2017)的方法与 CFMF-Net1 具有相似的训练集和识别网络,将它们单独对比。可以看到,CFMF-Net1 取得了更好的识别效果。因为 Masi 等人(2017)提出的是基于 3D 模型规则进行正面化的方法,生成的正面人脸不够逼真,Luan 等人(2017)和 Yin 等人(2017)提出的 2D 回归生成方法没有充分保留原图中的有效信息。而 CFMF-Net1 结合了 3D 和 2D 方法的优势,既保持了原始身份信息,又保证了生成图像足够逼真。CFMF-Net1 在 LFW 和 IJB-A 数据集上的正面化结果示例分别如图 6 和图 7 所示。本文方法 CFMF-Net2 是仅通过简单的形变场回归来正面化人脸,与结合了 GAN 与密集形变场的方法(Zhao 等,2018a;Cao 等,2018)相比,得到了与这些复杂方法持平的效果。

表 6 不同方法在 IJB-A 数据集上验证和识别准确率
Table 6 Face verification and identification accuracy of different methods on IJB-A dataset

方法	类型	人脸验证		人脸识别		训练集 图像转化/识别	网络结构
		FAR = 0.01	FAR = 0.001	Top-1	Top-5		
Crosswhite 等人(2017)	特征解耦	93.9	83.6	92.8	97.7	- / VGGFace	VGGNet16
Yang 等人(2017)	特征解耦	94.1	88.1	95.8	98.0	- / 3M ims of 50K ids	GoogLeNet-BN
Zhao 等人(2017)	特征解耦	97.6	93.0	97.1	98.9	300 W-LP/MS-Celeb	ResNeXT50 + GoogLeNet-BN
Zhu 等人(2016)	人脸增广	89.0	82.8	90.3	92.8	300 W-LP/MS-Celeb	-
Masi 等人(2017)	人脸增广	88.8	75.0	92.6	96.6	- / Webface	VGGNet
Chang 等人(2017)	人脸增广	90.1	85.2	91.4	93.0	- / MS-Celeb	ResNet101
Luan 等人(2017)	人脸正面化	77.4	53.9	85.5	95.7	MultiPIE/ Webface	CASIA-Net
Yin 等人(2017)	人脸正面化	85.2	66.3	90.2	95.4	300 W-LP/ Webface	CASIA-Net
Zhao 等人(2018a)	人脸正面化	93.3	87.5	94.4	-	MultiPIE/ -	LightCNN-29
Cao 等人(2018)	人脸正面化	95.2	89.7	96.1	97.9	CelebA-HQ/ -	LightCNN-29
CFMF-Net1(本文)	人脸正面化	90.7	77.3	94.7	98.1	300 W-LP/ Webface	Fast AlexNet
CFMF-Net2(本文)	人脸正面化	95.3	86.4	95.4	98.4	300 W-LP/ Webface	LightCNN-29

注:加粗字体表示各列最优结果,“-”表示原方法未报告结果。



图 6 CFMF-Net1 在 LFW 上的正面化结果示例
Fig. 6 Exemplars of frontalization results on LFW of CFMF-Net1 ((a) original input images; (b) x-axis morphing field; (c) frontalized results)



图 7 CFMF-Net1 在 IJB-A 上的正面化结果示例
Fig. 7 Exemplars of frontalization results on IJB-A of CFMF-Net1 ((a) original input images; (b) x-axis morphing field; (c) frontalized results)

值得一提的是,当前数据集的人脸图像主要的变化在 yaw 方向,即本文中的 x 方向。一种自然的想法是能否通过加强 x 方向形变场的训练权重来提升性能。然而实际上这种做法对性能几乎没有影响,因为 CFMF-Net 可以自动学习到形变场的主要变化在 x 方向。此外,给 x 方向形变场更多训练权重可能对可扩展性有影响,因为现实中的人脸图像还会存在其他方向上的姿态变化。

3.3 消融实验

为了分析 CFMF-Net 每个模块对人脸正面化和识别的影响,进行了一系列消融实验。在 300 W-LP 数据集上消融实验的可视化结果如图 8 所示。可以看到,通过 TPS 可以得到一个基本的人脸正面化结果(图 8(b))。直接利用粗粒度形变场得到的人脸正面化图像,由于自遮挡问题,依然存在一定程度的失真(图 8(c))。而借助细粒度形变场,可以得到逼真的正面化人脸图像(图 8(d))。这验证了 CFMF-Net 各部分对正面化的作用。

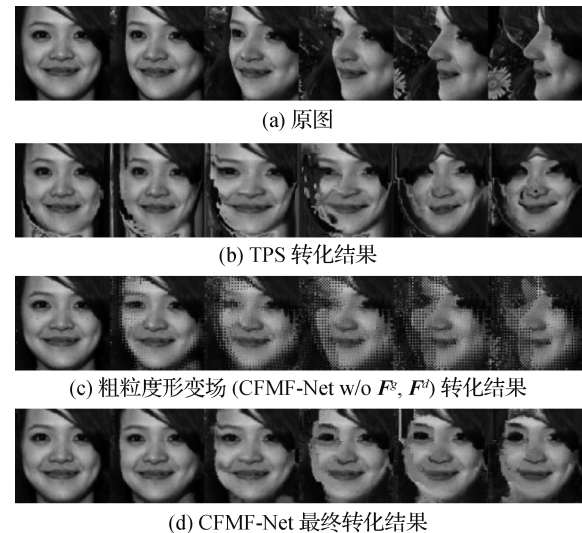


图 8 CFMF-Net 在 300 W-LP 上消融实验的结果
Fig. 8 Ablation study of frontalization on 300 W-LP ((a) original input images; (b) results of TPS; (c) results of CFMF-Net w/o F^g, F^d ; (d) results of CFMF-Net)

从识别结果的角度来看,CFMF-Net 的每一部分对人脸识别的准确率都具有重要作用。CFMF-Net 在 IJB-A 数据集上的消融实验结果如表 7 所示。可以看出,相比于不进行人脸正面化直接用 Fast Alex-Net 进行人脸识别,使用粗粒度形变场进行正面化,

表 7 CFMF-Net 在 IJB-A 上的消融实验
Table 7 Ablation study of CFMF-Net on IJB-A

方法	人脸验证准确率		人脸识别准确率	
	FAR = 0.01	FAR = 0.001	Top-1	Top-5
w/o CFMF-Net	85.2	66.0	90.9	97.1
CFMF-Net w/o F^g, F^d	88.8	70.7	92.1	97.2
CFMF-Net	90.7	77.3	94.7	98.1

注:加粗字体表示各列最优结果。

能在一定程度上提升人脸识别的准确率。而使用细粒度形变场进行人脸正面化,能进一步提升识别的准确率。

为了进一步验证 CFMF-Net 对大姿态人脸的效果,将 IJB-A 测试集按姿态大小分为 3 组,即 $[0^\circ, \pm 30^\circ)$ 、 $[\pm 30^\circ, \pm 60^\circ)$ 和 $[\pm 60^\circ, \pm 90^\circ)$ (详见<https://github.com/whobefore/MF-Net/tree/master/Data/IJBA>)。测试协议与 IJB-A 人脸识别测试相同,但每组再细分为 3 组不同姿态的实验,即 $[0^\circ, \pm 30^\circ)$ 的子集作为 gallery, $[0^\circ, \pm 30^\circ)$ 、 $[\pm 30^\circ, \pm 60^\circ)$ 、 $[\pm 60^\circ, \pm 90^\circ)$ 作为 probe 分别进行人脸识别测试。在每组数据上,首先用 CFMF-Net 进行人脸正面化,再用 Fast AlexNet 进行人脸识别,以测试识别准确率,并将其与直接使用 Fast AlexNet 进行识别的准确率相比较,结果如表 8 所示。可以看出,本文方法相比通用人脸识别方法(Deng 等,2019),在能力相当的网络结构下取得了更好结果,说明现在仍存在对姿态特殊处理的必要。另外,在大姿态 $[\pm 60^\circ, \pm 90^\circ)$ 的测试集上,正面化后图像的识别率得到显著提升,进一步验证了本文方法对大姿态人脸识别的有效性。

表 8 IJB-A 上不同姿态子集的 TOP-1 识别率
Table 8 Top-1 recognition accuracy in our self-defined pose-subdivision test protocol on IJB-A

方法	$[0^\circ, \pm 30^\circ)$	$[\pm 30^\circ, \pm 60^\circ)$	$[\pm 60^\circ, \pm 90^\circ)$
w/o CFMF-Net	96.2	89.1	73.9
Deng 等人(2019)	97.6	95.2	85.2
本文	98.8	97.1	87.6

注:加粗字体表示各列最优结果。

4 结 论

针对大姿态人脸识别问题,本文提出了一种基于由粗到细形变场回归的人脸正面化的方法 CFMF-Net。在实验结果中,尤其是大姿态的人脸识别实验中,本文方法表现出了比相关方法更好或持平的效果,表明该方法可以有效结合 2D 和 3D 人脸正面化方法的优点,既充分保留了原始图像中的信息,又保证了生成的正面图像足够逼真。与通用人脸识别方法的对比结果表明,尽管可以通过数据集的丰富和

损失函数的设计显著提升直接进行人脸识别方法的性能,但目前对人脸姿态的处理仍然存在其必要性。然而在本文方法中,虽然通过由粗到细的学习方式提升了密集形变场回归的鲁棒性,但这样的算法仍然有很高的自由度,压缩形变场的冗余信息是一种更好的解决方式。在未来的工作中,一方面希望对密集形变场进行结构可保持的稀疏化,另一方面希望能够进一步设计出识别性能驱动的自动人脸或人脸特征对齐方法,发掘出最佳人脸对齐角度,并应用到更复杂场景的人脸识别中。

参考文献 (References)

Asthana A, Marks T K, Jones M J, Tieu K H and Rohith M V. 2011. Fully automatic pose-invariant face recognition via 3D pose normalization//Proceedings of 2011 International Conference on Computer Vision. Barcelona, Spain: IEEE: 937-944 [DOI: 10.1109/ICCV.2011.6126336]

Bookstein F L. 1989. Principal warps: thin-plate splines and the decomposition of deformations. IEEE Transactions on Pattern Analysis and Machine Intelligence, 11 (6): 567-585 [DOI: 10.1109/34.24792]

Cao J, Hu Y B, Zhang H W, He R and Sun Z N. 2018. Learning a high fidelity pose invariant model for high-resolution face frontalization//Proceedings of the 32nd International Conference on Neural Information Processing Systems. Montréal, Canada: Curran Associates Inc.: 2872-2882

Chang F J, Tran A T, Hassner T, Masi I, Nevatia R and Medioni G. 2017. FacePoseNet: making a case for landmark-free face alignment//Proceedings of 2017 IEEE International Conference on Computer Vision Workshops. Venice, Italy: IEEE: 1599-1608 [DOI: 10.1109/ICCVW.2017.188]

Crosswhite N, Byrne J, Stauffer C, Parkhi O, Cao Q and Zisserman A. 2017. Template adaptation for face verification and identification//Proceedings of the 12th IEEE International Conference on Automatic Face and Gesture Recognition. Washington, USA: IEEE: 1-8 [DOI: 10.1109/FG.2017.11]

Deng J K, Guo J, Xue N N and Zafeiriou S. 2019. ArcFace: additive angular margin loss for deep face recognition//Proceedings of 2019 IEEE/CVF Conference on Computer Vision and Pattern Recognition. Long Beach, USA: IEEE: 4685-4694 [DOI: 10.1109/CVPR.2019.00482]

Ding C X, Xu C and Tao D C. 2015. Multi-task pose-invariant face recognition. IEEE Transactions on Image Processing, 24(3): 980-993 [DOI: 10.1109/TIP.2015.2390959]

Goodfellow I J, Pouget-Abadie J, Mirza M, Xu B, Warde-Farley D, Ozair S, Courville A and Bengio Y. 2014. Generative adversarial

- nets//Proceedings of the 27th International Conference on Neural Information Processing Systems. Montreal, Canada; MIT Press; 2672-2680
- Hu L Q, Kan M N, Shan S G, Song X G and Chen X L. 2017. LDF-Net: learning a displacement field network for face recognition across pose//Proceedings of the 12th IEEE International Conference on Automatic Face and Gesture Recognition. Washington, USA; IEEE; 9-16 [DOI: 10.1109/FG.2017.12]
- Huang G B and Learned-Miller E. 2014. Labeled Faces in the Wild: Updates and New Reporting Procedures. Amherst Technical Report UM-CS-2014-003. University of Massachusetts
- Huang R, Zhang S, Li T Y and He R. 2017. Beyond face rotation: global and local perception GAN for photorealistic and identity preserving frontal view synthesis//Proceedings of 2017 IEEE International Conference on Computer Vision. Venice, Italy; IEEE; 2458-2467 [DOI: 10.1109/ICCV.2017.267]
- Kan M N, Shan S G, Chang H and Chen X L. 2014. Stacked progressive auto-encoders (SPA-E) for face recognition across poses//Proceedings of 2014 IEEE Conference on Computer Vision and Pattern Recognition. Columbus, USA; IEEE; 1883-1890 [DOI: 10.1109/CVPR.2014.243]
- Kan M N, Shan S G, Zhang H H, Lao S H and Chen X L. 2016. Multi-view discriminant analysis. IEEE Transactions on Pattern Analysis and Machine Intelligence, 38(1): 188-194 [DOI: 10.1109/TPAMI.2015.2435740]
- Klare B F, Klein B, Taborsky E, Blanton A, Cheney J, Allen K, Grother P, Mah A, Burge M and Jain A K. 2015. Pushing the frontiers of unconstrained face detection and recognition: IARPA Janus benchmark A//Proceedings of 2015 IEEE Conference on Computer Vision and Pattern Recognition. Boston, USA; IEEE; 1931-1939 [DOI: 10.1109/CVPR.2015.7298803]
- Li A N, Shan S G, Chen X L and Gao W. 2009. Maximizing intra-individual correlations for face recognition across pose differences//Proceedings of 2009 IEEE Conference on Computer Vision and Pattern Recognition. Miami, USA; IEEE; 605-611 [DOI: 10.1109/CVPR.2009.5206659]
- Li S X, Liu X, Chai X J, Zhang H H, Lao S H and Shan S G. 2012. Morphable displacement field based image matching for face recognition across pose//Proceedings of the 12th European Conference on Computer Vision. Florence, Italy; Springer; 102-115 [DOI: 10.1007/978-3-642-33718-5_8]
- Luan T, Yin X and Liu X M. 2017. Disentangled representation learning GAN for pose-invariant face recognition//Proceedings of 2017 IEEE Conference on Computer Vision and Pattern Recognition. Honolulu, USA; IEEE; 1283-1292 [DOI: 10.1109/CVPR.2017.141]
- Luan X, Geng H M, Liu L H, Li W S, Zhao Y Y and Ren M. 2020. Geometry structure preserving based GAN for multi-pose face frontalization and recognition. IEEE Access, 8: 104676-104687 [DOI: 10.1109/ACCESS.2020.2996637]
- Masi I, Hassner T, Tràn A T and Medioni G. 2017. Rapid synthesis of massive face sets for improved face recognition//Proceedings of the 12th IEEE International Conference on Automatic Face and Gesture Recognition. Washington, USA; IEEE; 604-611 [DOI: 10.1109/FG.2017.76]
- Peng X, Yu X, Sohn K, Metaxas D N and Chandraker M. 2017. Reconstruction-based disentanglement for pose-invariant face recognition//Proceedings of 2017 IEEE Conference on Computer Vision and Pattern Recognition. Venice, Italy; IEEE; 1632-1641 [DOI: 10.1109/ICCV.2017.180]
- Prabhu U, Heo J and Savvides M. 2011. Unconstrained pose-invariant face recognition using 3D generic elastic models. IEEE Transactions on Pattern Analysis and Machine Intelligence, 33(10): 1952-1961 [DOI: 10.1109/TPAMI.2011.123]
- Rong C L, Zhang X M and Lin Y B. 2020. Feature-improving generative adversarial network for face frontalization. IEEE Access, 8: 68842-68851 [DOI: 10.1109/ACCESS.2020.2986079]
- Sagonas C, Antonakos E, Tzimiropoulos G, Zafeiriou S and Pantic M. 2016. 300 faces in-the-wild challenge: database and results. Image and Vision Computing, 47: 3-18 [DOI: 10.1016/j.imavis.2016.01.002]
- Sengupta S, Chen J C, Castillo C, Patel V M, Chellappa R and Jacobs D W. 2016. Frontal to profile face verification in the wild//Proceedings of 2016 IEEE Winter Conference on Applications of Computer Vision (WACV). Lake Placid, USA; IEEE; 1-9 [DOI: 10.1109/WACV.2016.7477558]
- Sharma A and Jacobs D W. 2011. Bypassing synthesis: PLS for face recognition with pose, low-resolution and sketch//Proceedings of 2011 CVPR. Colorado Springs, USA; IEEE; 593-600 [DOI: 10.1109/CVPR.2011.5995350]
- Sim T, Baker S and Bsat M. 2003. The CMU pose, illumination, and expression database. IEEE Transactions on Pattern Analysis and Machine Intelligence, 25(12): 1615-1618 [DOI: 10.1109/TPAMI.2003.1251154]
- Wu X, He R, Sun Z N and Tan T N. 2018. A light CNN for deep face representation with noisy labels. IEEE Transactions on Information Forensics and Security, 13(11): 2884-2896 [DOI: 10.1109/TIFS.2018.2833032]
- Yang J L, Ren P R, Zhang D Q, Chen D, Wen F, Li H D and Hua G. 2017. Neural aggregation network for video face recognition//Proceedings of 2017 IEEE Conference on Computer Vision and Pattern Recognition. Honolulu, USA; IEEE; 5216-5225 [DOI: 10.1109/CVPR.2017.554]
- Yi D, Lei Z, Liao S C and Li S Z. 2014. Learning face representation from scratch [EB/OL]. [2021-01-18]. <https://arxiv.org/pdf/1411.7923.pdf>
- Yin X and Liu X M. 2018. Multi-task convolutional neural network for pose-invariant face recognition. IEEE Transactions on Image Processing, 27(2): 964-975 [DOI: 10.1109/TIP.2017.2765830]

- Yin X, Xiang Y, Sohn K, Liu X M and Chandraker M. 2017. Towards large-pose face frontalization in the wild//Proceedings of 2017 IEEE International Conference on Computer Vision. Venice, Italy: IEEE: 4010-4019 [DOI: 10.1109/ICCV.2017.430]
- Zhang S F, Miao Q H, Huang M, Zhu X Y, Chen Y Y, Lei Z and Wang J Q. 2019. Pose-weighted GAN for photorealistic face frontalization//Proceedings of 2019 IEEE International Conference on Image Processing. Taipei, China: IEEE: 2384-2388 [DOI: 10.1109/ICIP.2019.8803362]
- Zhang Y Z, Shao M, Wong E K and Fu Y. 2013. Random faces guided sparse many-to-one encoder for pose-invariant face recognition//Proceedings of 2013 IEEE International Conference on Computer Vision. Sydney, Australia: IEEE: 2416-2423 [DOI: 10.1109/ICCV.2013.300]
- Zhao J, Cheng Y, Xu Y, Xiong L, Li J S, Zhao F, Jayashree K, Pranata S, Shen S M, Xing J L, Yan S C and Feng J S. 2018a. Towards pose invariant face recognition in the wild//Proceedings of 2018 IEEE/CVF Conference on Computer Vision and Pattern Recognition. Salt Lake City, USA: IEEE: 2207-2216 [DOI: 10.1109/CVPR.2018.00235]
- Zhao J, Xiong L, Cheng Y, Cheng Y, Li J S, Zhou L, Xu Y, Karlekar J, Pranata S, Shen S M, Xing J L, Yan S C and Feng J S. 2018b. 3D-aided deep pose-invariant face recognition//Proceedings of the 27th International Joint Conference on Artificial Intelligence. Stockholm, Sweden: IJCAI: 1184-1190 [DOI: 10.24963/ijcai.2018/165]
- Zhao J, Xiong L, Jayashree K, Li J S, Zhao F, Wang Z C, Pranata S, Shen S M, Yan S C and Feng J S. 2017. Dual-agent GANs for photorealistic and identity preserving profile face synthesis//Proceedings of the 31st International Conference on Neural Information Processing Systems. Long Beach, USA: Curran Associates Inc.: 65-75
- Zhu X Y, Lei Z, Liu X M, Shi H L and Li S Z. 2016. Face alignment across large poses: a 3D solution//Proceedings of 2016 IEEE Conference on Computer Vision and Pattern Recognition. Las Vegas, USA: IEEE: 146-155 [DOI: 10.1109/CVPR.2016.23]

- Zhu X Y, Lei Z, Yan J J, Yi D and Li S Z. 2015. High-fidelity pose and expression normalization for face recognition in the wild//Proceedings of 2015 IEEE Conference on Computer Vision and Pattern Recognition. Boston, USA: IEEE: 787-796 [DOI: 10.1109/CVPR.2015.7298679]
- Zhu Z Y, Luo P, Wang X G and Tang X O. 2013. Deep learning identity-preserving face space//Proceedings of 2013 IEEE International Conference on Computer Vision. Sydney, Australia: IEEE: 113-120 [DOI: 10.1109/ICCV.2013.21]

作者简介



胡蓝青, 1992年生, 女, 博士研究生, 主要研究方向为计算机视觉、人脸识别和迁移学习。

E-mail: lanqing.hu@vip.ict.ac.cn



山世光, 通信作者, 男, 研究员, 主要研究方向为计算机视觉、模式识别和机器学习。

E-mail: sgshan@ict.ac.cn

阚美娜, 女, 副研究员, 主要研究方向为计算机视觉、模式识别、迁移学习和弱监督学习。E-mail: kanmeina@ict.ac.cn

陈熙霖, 男, 研究员, 主要研究方向为计算机视觉、模式识别、多媒体技术和多模式人机接口。E-mail: xlchen@ict.ac.cn