

基于 GAN 网络的三维人脸纹理重建技术

程慧敏¹ 姚剑敏^{1,2} 陈恩果¹ 严群^{1,2}

CHENG Huimin YAO Jianmin CHEN En'guo YAN Qun

摘要

在计算机视觉领域,三维人脸重建起着重要的作用,特别是在动画制作、虚拟人脸、人脸识别等领域正被广泛应用。基于单张图片的三维人脸重建,通过使用拟合的模型对图像进行采样,可以创建出面部的UV纹理图。然而通过相机采集到的是单视角的二维图像,存在人脸自遮挡的情况,这导致生成的UV纹理图的信息是不完整的。因此,提出一种针对三维人脸纹理补全问题的条件生成对抗网络,将Unet和部分卷积结合起来作为生成器,从而在纹理修复时可以保留更多的图像信息;判别器中引入SM-PatchGAN提高图像判别的真实准确度。实验结果表明,相较于其他算法,所提出的算法取得了较为显著的改进效果,特别是在处于大视角下存在人脸自遮挡问题时,也能重建出精细的三维人脸模型。

关键词

生成对抗网络;深度学习;三维人脸重建;纹理补全;自遮挡

doi: 10.3969/j.issn.1672-9528.2024.03.051

0 引言

近年来,虚拟现实、虚拟人脸生成技术逐渐走进人们的视野,相应的三维人体和三维人脸重建技术也随着社会科技的发展在不断进步,其需求也越来越多。而如何生成更为真实的三维人脸模型,一直是目前要解决的一个重要课题。从单张二维图像中重建出完整的三维人脸模型常常会面临相机拍摄角度偏移或者人脸被物体遮挡导致人脸面部纹理信息缺失的问题^[1],这会导致重建的三维人脸模型在遮挡区域的纹理信息缺失或不准确。

在过去的几年里,人们见证了在稀疏和密集的三维面对齐方面取得的相当大的进展。其中一些发展包括使用深度神经网络来恢复三维面部结构^[2-4],同时在三维人脸重建方面也取得了许多成果。如今三维人脸重建主要围绕三维人脸可形变模型(3D morphable model, 3DMM)展开,例如Blanz等人在文献[5]和文献[6]中提出的3DMM方法,输入图像可以是单张或多张的正面人脸图像,也可以是用户交互输入的图像。近年来一些方法(如文献[2]和文献[7])将卷积神经网络(CNN)引入到3DMM模型中,使用卷积神经网络对3DMM模型的参数进行直接预测。但基于深度学习的三维人脸重建方法大部分都属于有监督学习,需要大量的标注数据,

然而现有的真实三维人脸数据较少,需要花费大量时间标注。因此,Deng等人^[8]提出了一种弱监督的方法,实现三维人脸重建。

现有的三维人脸重建方法大多都基于正面的人脸图像,而面对大视角下出现的自遮挡和存在物体遮挡面部的情况时,重建的面部模型存在纹理信息缺失的问题。针对人脸视角偏移导致纹理缺失的问题,Pathak等人^[9]提出了具有重建和对抗性损失的上下文编码器,以生成符合邻近区域的缺失区域的内容。文献[10]中的方法结合了一个重构损失、两个对抗性损失和一个语义解析损失,以确保局部-全局内容的真实性 and 一致性。然而,他们的方法只在一小部分预定义的掩模上进行训练,并不能直接解决所面临的问题。文献[11]学习在不使用任何完整纹理的情况下完成单个人脸图像中的不可见纹理,提出DSD-GAN,利用不同主题的人脸图像数据集,以一种无监督的方式训练一个纹理补全模型。文献[12]在采用条件生成对抗网络(conditional generative adversarial networks, CGAN)^[13]的基础上,使用了全局判别器对整体面部纹理判别,并提出了局部判别器,以提升面部中心区域纹理的细节,但是在面对存在大角度的人脸自遮挡情况时,其补全结果与原图仍有较大差异。

1 提出的方法

1.1 本文算法

本文算法结构如图1所示,首先根据一个残差网络训练一个回归网络,拟合出3DMM模型。在得到3DMM模型后,从中提取出缺失部分纹理的UV纹理图。训练了一个基于条

1. 福州大学 福建福州 350108

2. 晋江市博感电子科技有限公司 福建晋江 362200

[基金项目] 国家重点研发计划(2022YFB3603503);福建省技术攻关重点项目(2023G007)

件生成对抗网络作为纹理补全网络，本文提出的纹理补全网络，加入部分卷积层^[14]的U-net^[15]网络作为生成器，同时使用SM-PatchGAN（soft mask-guided PatchGAN）^[16]作为全局判别器，使用其中的生成器来补全缺失的UV纹理图。

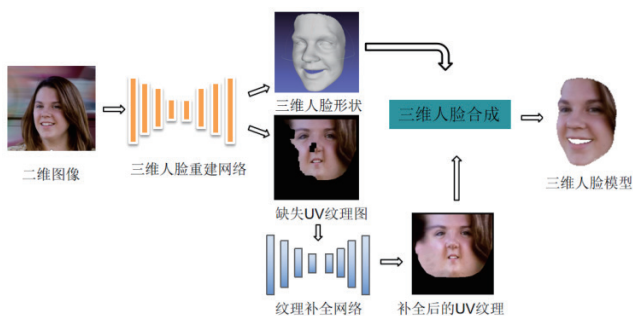


图1 三维人脸纹理补全算法流程图

1.2 3DMM 模型

为了从二维图像中重建出三维人脸模型，采用了一种弱监督的方法拟合出3DMM模型。3DMM的核心思想是一个3维的人脸，可以由其他人脸线性组合。每一张人脸可以由形状向量 S 和纹理向量 T 组成，可以表示为：

$$S = S(\alpha, \beta) = \bar{S} + B_{id}\alpha + B_{exp}\beta \quad (1)$$

$$T = T(\delta) = \bar{T} + B_t\delta \quad (2)$$

式中： $\bar{S}$ 和 $\bar{T}$ 分别表示人脸的平均形状和平均纹理， B_{id} 、 B_{exp} 和 B_t 表示人脸身份、表情和纹理的PCA基，而 α 、 β 和 δ 分别为对应的人脸身份、表情和纹理的PCA系数。人脸模型采用了2009 Basel Face Model^[17]中的 $\bar{S}$ 、 $\bar{T}$ 、 B_{id} 和 B_t 以及Guo等人^[18]中的 B_{exp} 。使用球谐函数（spherical harmonics, SH）估计场景光照。相机模型使用透视投影。

1.3 UV 纹理提取

人脸UV纹理提取模块的目的是在3DMM模型中提取出对应的UV纹理。在3DMM模型中纹理通常用 $T \in \mathbb{R}^{3 \times N}$ 的向量表征，这种方式会导致人脸纹理丢失了空间结构性。因此在本文中，采用UV图来表示人脸纹理。UV贴图也被称为纹理映射，它将二维图像与三维模型表面上的每个顶点相关联，然后将三维人脸纹理信息映射到二维空间的坐标系下，它可以储存大量的细节和信息，如皮肤斑点、皱纹等。同时，在三维人脸重建任务中，采用UV贴图来表示人脸纹理，能更方便地将其送入生成对抗网络中进行训练，生成高质量的补全纹理。

假设三维模型的每个顶点坐标表示为 $v_i=(x, y, z)$ ，其对应的UV空间投影为 $v_i=(u, v)$ ，可以得到三维人脸顶点坐标与UV空间坐标的对应关系为：

$$u = a_2 y + b_2 \quad (3)$$

$$v = a_1 \tan^{-1}\left(\frac{x}{z}\right) + b_1 \quad (4)$$

式中： a_1 、 a_2 是恒定尺度因子， b_1 、 b_2 为平移系数，用于将人脸展开到图像边界中。

将对应顶点映射到UV空间后，为获得精细的三维人脸纹理信息，需再将三维顶点的RGB三通道色彩信息赋值到UV空间中，即可得到对应人脸的UV贴图，它包含着所需要的纹理信息。但由于自遮挡的存在，此时生成的UV纹理贴图是残缺纹理，存在部分纹理不可见的现象，要依靠后续的纹理补全网络才能得到完整的三维人脸模型。

为了判断每个点是否被遮挡，采用z-buffer算法，使用双线性采样器得到投影到人脸图片上的投影点的Z缓冲深度 $z_buffer(t_x, t_y)$ 。然后将其对应的空间三维顶点的深度值 t_z 与Z缓冲深度进行比较，若 $t_z > z_buffer(t_x, t_y)$ ，就视为该点被遮挡。由此可以判断哪些顶点是可见的，哪些顶点是不可见的，从而利用这个Z缓冲生成关于UV纹理图的可见性Mask。

1.4 网络结构

1.4.1 生成器结构

现有的面部纹理补全方法假设缺失的区域是规则的、矩形的。文献[12]提出的UV-GAN网络应用了局部判别器来区分生成的内容和真实的内容，将填充好的中心缺失区域输入到局部判别器中，增强局部纹理细节。但是当自遮挡造成的人脸面部纹理缺失时，由于人脸面对相机时角度的不同，导致面部纹理的缺失区域通常是不规则的。因此，本文提出了一种采用部分卷积层代替U-net中标准卷积层的方法。

U-net是一种全卷积神经网络，它分为编码器和解码器两个部分。编码器将图像进行降采样，同时增加特征图的通道数；解码器通过上采样的方式将图像逐渐恢复，同时转换为完整的纹理补全结果。U-net输出的结果是与输入图像尺寸一致的高清图像。将U-net和部分卷积结合起来，则是通过在U-net解码器的每一层上增加部分卷积层，用于修复图像中的缺失纹理区域。传统的卷积层一般对未遮挡区域像素和Mask像素无差别地做卷积，这会导致纹理的缺失或是边缘出现伪影。与传统的卷积层不同，部分卷积只对已知部分的像素进行卷积运算，保留遮挡部分原本的像素值，不对其进行运算，从而在遮挡修复时可以保留更多的图像信息。部分卷积的计算公式为：

$$x' = \begin{cases} W^T(X \otimes M) \frac{\text{sum}(1)}{\text{sum}(M)} + b, & \text{sum}(M) > 0 \\ 0, & \text{otherwise} \end{cases} \quad (5)$$

式中： $\otimes$ 表示元素的乘法，1是全1的矩阵，其与 M 的形状相同。可以看出，部分卷积的结果仅由有效的输入值（如

$X \otimes M$) 决定。 $sum(1)/sum(M)$ 是一个比例因子, 当部分卷积的输入值发生数量上的变化时, 它会调整对应结果。

在每个部分卷积层之后更新二进制掩码, 更新二进制掩码的规则如公式 (6) 所示。将有效区域的总和与有效区域的像素数量进行比较, 如果总和小于像素数量, 说明输入数据的有效区域部分受到遮挡。经过部分卷积层后, 将对应位置的二进制掩码更新, 更新的二进制 Mask 中有更少的 0。随着部分卷积的次数不断增加, 二进制 Mask 最终将全部更新为 1。这意味着不需要考虑缺失区域的大小和形状。

$$m' = \begin{cases} 1, & sum(M) > 0 \\ 0, & otherwise \end{cases} \quad (6)$$

1.4.2 判别器结构

SM-PatchGAN 的工作过程主要分为两部分: 首先, 先把补全得到纹理图像切成 $N \times N$ 的小块 (Patch); 其次, SM-PatchGAN 将产生的 $N \times N$ 的小块分组输入到多个小型判别器中判别和分类, 根据图像的几何形状和像素差异输出结果。它的输出结果是每个小块的置信度, 而不是整个图像的置信度; 最终, 将每个小块的置信度值求平均, 得到整个图像的置信度。一张由 GAN 网络补全得到的 UV 纹理图包含未遮挡的真实纹理部分和由生成器生成的补全纹理部分。传统的 PatchGAN 判别器是将真实纹理图的判别输出的特征图拉向“真值 1”, 而判断为生成器生成的纹理图的输出特征图全部判为“假值 0”, 这忽略了未遮挡区域确实是真实图像。SM-PatchGAN 的改进在于将判别输出中由生成器生成的纹理部分拉向“假值 0”, 而将判别输出中来自未遮挡区域的真实纹理部分拉向“真值 1”。但是, 这样处理虽然提升了补全部分的生成图像的质量, 但又使得生成的纹理部分与真实纹理部分的特征融合不够充分, 并且部分纹理图像的边界周围可能同时包含真实像素和合成像素, 因此, SM-PatchGAN 又引入高斯滤波进行模糊处理来解决这个问题, 增强判别器的能力。

1.4.3 整体结构设计

本文采用条件生成对抗网络进行 UV 纹理补全。该生成对抗网络由一个引入部分卷积的 U-net 作为生成器以及一个改进的 PatchGAN 作为判别器。生成器采用 U-net 网络结构, 融合底层细粒度特征和高层抽象; 判别器采用 PatchGAN 网络结构, 在图块尺度提取纹理等高频信息, 同时为了优化纹理补全部分, 采用 SM-PatchGAN 能够将缺失区域的合成部分与上下文的真实部分区分开来, 从而增强判别器的能力。将 UV 残缺纹理图代替噪声输入生成器中, 并将 UV 完整纹理图和生成的 UV 纹理图放入判别器中训练。改进的网络结

构图如图 2 所示。

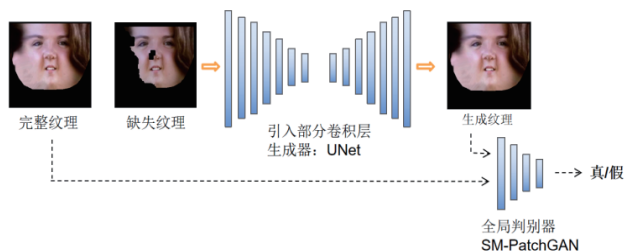


图 2 网络结构图

1.5 损失函数设计

为了保证 UV 纹理的细节和完整度, 采用以下几种损失函数的加权和作为最终的损失函数。

1.5.1 像素层损失函数 L_{rec}

采用像素级的 L_{rec} 范数作为重建损失:

$$L_{rec} = \|G(x) - y\| \quad (7)$$

式中: $G(x)$ 、 x 与 y 分别为生成的 UV 纹理、残缺 UV 纹理和完整 UV 纹理。

1.5.2 对称损失函数 L_{sym}

人脸纹理存在左右对称的特性, 因此可以利用对称来消除由自遮挡引起的纹理缺失问题。对称损失函数 L_{sym} 利用可见部分像素, 来补全缺失部分的纹理。并且仅约束不可见部分, 保留可见像素。对抗损失函数定义为:

$$L_{sym} = \|(G(x) - y_{flip}) \times m\| \quad (8)$$

将 m 设置为缺失区域 Mask, 其中可见像素的值为 0, y_{flip} 为对称部分的真实 UV 纹理。

1.5.3 对抗损失函数 L_{adv}^D 、 L_{adv}^G

对抗性损失反映了发生器如何最大限度地欺骗判别器, 以及判别器如何区分真假的能力, 定义为:

$$L_{adv}^D = (G, D) = E_{x \sim p_x} [(D(G(x)) - \sigma(1-m))^2] + E_{y \sim p_{data}} [(D(y) - 1)^2] \quad (9)$$

$$L_{adv}^G = E_{x \sim p_x} [(D(G(x)) - 1)^2 \times m] \quad (10)$$

式中: p_x 和 p_{data} 表示输入数据和真实数据的分布; σ 为降采样和高斯滤波的组合函数。只将预测得到的缺失区域的合成图像部分来优化生成器, 通过这种优化, 提升判别器判别出残缺纹理区域的合成块的能力。增强的判别器反过来可以促进生成器合成更真实的纹理。将最终的整体网络损失函数定义为:

$$L = L_{rec} + \lambda_1 L_{sym} + \lambda_2 L_{adv}^D + \lambda_3 L_{adv}^G \quad (11)$$

式中: λ_1 、 λ_2 、 λ_3 用来平衡各损失函数。

2 实验与结果与分析

2.1 数据集

本次实验所采用的数据集是 300W_LP^[19] 数据集。300W_LP 主要包含 8 个数据库: AFW、AFW_Flip、HELEN、

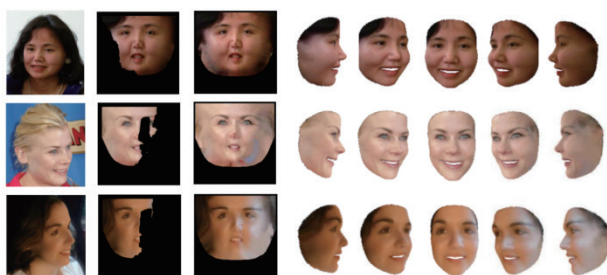
HELEN_Flip、IBUG、IBUG_Flip、LFPW、LFPW_Flip。300W_LP 共有数据 122 450 张图片。每张图片对应的标签都以 mat 形式保存，标签内容包括面部形状、表情、颜色、姿态参数等 3DMM 参数。将 300W_LP 数据集的正面人脸图片输入到本文使用的三维人脸重建网络，回归得到对应的三维人脸模型，然后从中提取出完整面部 UV 纹理图。每张图片在 Z 轴分别正向逆向每间隔 10° 旋转一次，将该正面人脸图片对应的 -90° ~ 90° 范围内的旋转图同样输入三维人脸重建网络，获得大量面部 UV 残缺纹理图。将完整面部 UV 纹理图作为残缺纹理图的标签。UV 纹理图分辨率为 256×256，完整 UV 纹理图作为判别器的输入。残缺纹理图代替噪声作为 GAN 的输入进行训练。

2.2 实验条件

实验在 Ubuntu 18.04、Window 10 环境下进行，使用的编程环境为 PyCharm，编程语言使用的是 Python 3.7，深度学习框架使用的是 Tensorflow 1.14.0 以及 CUDA 10.1。硬件设备采用 i7-6700CPU+Tesla P40 显卡服务器，运行内存 16 GB (RAM)。

2.3 实验结果及分析

本文实验结果如图 3 所示。图 A 为存在大角度偏移的人脸输入图像。图 B 为由输入图像形成的 UV 残缺纹理图像，由于人脸存在大角度的自遮挡情况，将遮挡部分用黑色掩膜填充。图 C 为残缺纹理通过本文提出的 GAN 网络得到的补全纹理图。将补全得到的完整纹理图与三维人脸模型进行拟合，最终得到具有完整纹理的三维人脸模型。



A. 输入图像 B. 残缺纹理 C. 补全纹理 D. 五视角的三维人脸重建模型

图 3 带遮挡 UV 纹理、补全的 UV 纹理以及三维人脸拟合结果

本文采用峰值信噪比 (peak signal to noise ratio, PSNR) 和结构相似性指数 (structural similarity index, SSIM)^[20] 来评估补全纹理的图像质量。PSNR 用来直接测量像素值的差距，两个图像之间的 PSNR 值越大，则说明补全纹理与真实纹理越相似，补全效果越好。SSIM 是一种衡量两幅图像相似度的指标，同样，其值越大，代表相似性越高，纹理补全效果越好。PSNR 和 SSIM 表示为：

$$P = 10 \log_{10} \left(\frac{M_I^2}{E_{MSE}} \right) \quad (12)$$

$$M = \frac{1}{mn} \sum_{i=0}^{m-1} \sum_{j=0}^{n-1} [I(i, j) - K(i, j)]^2 \quad (13)$$

式中：P 为峰值信噪比 PSNR； E_{MSE} 为真实纹理图 I 与预测纹理图 K 的均方误差； M_I^2 为图片可能的最大像素值。

$$S(x, y) = [l(x, y)^\alpha \times c(x, y)^\beta \times s(x, y)^\gamma] \quad (14)$$

$$l(x, y) = \frac{2 \times \mu_x \times \mu_y + c_1}{\mu_x^2 + \mu_y^2 + c_1} \quad (15)$$

$$c(x, y) = \frac{2 \times \sigma_x \times \sigma_y + c_2}{\sigma_x^2 + \sigma_y^2 + c_2} \quad (16)$$

$$s(x, y) = \frac{\sigma_{xy} + c_3}{\sigma_x \times \sigma_y + c_3} \quad (17)$$

式中：S 为结构相似性指数 SSIM； $l(x, y)$ 、 $c(x, y)$ 、 $s(x, y)$ 分别为比较样本 x 与 y 的亮度、对比度和结构； μ_x 和 μ_y 分别表示 x 与 y 的平均值； σ_x 和 σ_y 分别表示 x 与 y 的标准差； σ_{xy} 为 x 和 y 的协方差； c_1 、 c_2 、 c_3 为常数。

其中 AOT-GAN、DSD-GAN、UV-GAN 结果直接来源于其论文的实验数据。AOT-GAN 在高质量名人人脸属性 (high quality celeb faces attribute, CelebA-HQ) 数据集上进行训练。DSD-GAN 是一种无监督的三维人脸纹理修复算法。UV-GAN 方法在 Multi-PIE 数据集上测试。对比结果如表 1 所示。

表 1 各模型算法 PSNR 与 SSIM 比较结果

算法模型	PSNR	SSIM
AOT-GAN	24.65	0.852
DSD-GAN	—	0.768
UV-GAN	27.5	0.912
Ours	28.99	0.918

通过对比结果可以看出，使用本文提出的网络架构生成补全纹理在 PSNR 和 SSIM 指标上都优于其他方法。这说明，本文提出的网络结构能够更好地修复缺失纹理，与完整纹理的相似度更高。

3 结语

本研究在 GAN 的生成器中加入了部分卷积，并使用 SM-PatchGAN 作为判别器，通过优化模型的结构和参数，实现了对缺失纹理的自动补全。在 300W-LP 数据集的基础上提取 UV 残缺纹理图和正面完整面部纹理图生成数据集，并进行实验验证，本文提出的方法能够有效还原缺失的纹理信息，同时保持纹理的真实性和准确性。

研究结果表明，使用部分卷积的生成器架构和 SM-PatchGAN 判别器可以提高三维人脸纹理补全的效果。本文的模型在 PSNR 和 SSIM 指标上分别达到了 28.99 和 0.918，明显优于其他模型。相较于传统方法，本文的方法在保证真实性的同时，能够更好地恢复纹理细节。此外，本文的方法

在不同角度的人脸自遮挡、不同程度的纹理缺失情况下,依然具有很好的鲁棒性和可扩展性。未来的研究将会尝试使用其他的网络结构来提高纹理补全的效果,实现自监督的三维人脸纹理补全,进一步提高生成效率。

参考文献:

- [1] SHANG J, SHEN T, LI S, et al. Self-supervised monocular 3D face reconstruction by occlusion-aware multi-view geometry consistency[EB/OL]. (2020-07-24). <https://arxiv.org/pdf/2007.12494.pdf>.
- [2] ELAD R, MATAN S, ROY OR-EL, et al. Learning detailed face reconstruction from a single image[C]//2017 IEEE Conference on Computer Vision and Pattern Recognition (CVPR). Piscataway: IEEE, 2017: 5553-5562.
- [3] KIM H, ZOLLHÖFER, MICHAEL, et al. InverseFaceNet: deep monocular inverse face rendering[C]//2018 IEEE/CVF Conference on Computer Vision and Pattern Recognition. Piscataway: IEEE, 2018: 4625-4634.
- [4] ZHU X Y, LEI Z, LIU X M, et al. Face alignment across large poses: a 3D solution[C]//29th IEEE Conference on Computer Vision and Pattern Recognition. Piscataway: IEEE, 2016: 146-155.
- [5] BLANZ V, VETTER T. A morphable model for the synthesis of 3D faces[M]//Seminal Graphics Papers: Pushing the Boundaries, 2023: 157-164.
- [6] BLANZ V, VETTER T. Face recognition based on fitting a 3D morphable model[J]. Pattern analysis and machine intelligence, IEEE transactions on, 2003, 25(9): 1063-1074.
- [7] DOU P, SHAH S K, KAKADIARIS I A. End-to-end 3D face reconstruction with deep neural networks[C]// IEEE/CVF Conference on Computer Vision and Pattern Recognition. Piscataway: IEEE, 2017: 1503-1512.
- [8] DENG Y, YANG J, XU S, et al. Accurate 3D face reconstruction with weakly-supervised learning: from single image to image set[C]//Proceedings Of The IEEE/Cvf Conference On Computer Vision And Pattern Recognition Workshops. Piscataway: IEEE, 2019: 285-295.
- [9] DEEPAK P, PHILIPP K, JEFF D, et al. Context encoders: feature learning by inpainting[C]// Conference on Computer Vision and Pattern Recognition (CVPR). Piscataway: IEEE, 2016: 2536-2544.
- [10] LI Y, LIU S, YANG J, et al. Generative face completion[EB/OL]. (2017-04-19)[2023-10-26]. <https://arxiv.org/pdf/1704.05838.pdf>.
- [11] KIM J Y, YANG J L, TONG X. Learning high-fidelity face texture completion without complete face texture[C]// IEEE international conference on computer vision (ICCV). Piscataway: IEEE, 2021: 13970-13979.
- [12] DENG J K, CHENG S Y, XUE N N, et al. UV-GAN: adversarial facial UV map completion for pose-invariant face recognition[C]// IEEE/CVF Conference on Computer Vision and Pattern Recognition. Piscataway: IEEE, 2017: 7093-7102.
- [13] PHILLIP I, JUNYAN Z, TINGHUI Z, et al. Image-to-image translation with conditional adversarial networks[C]// IEEE Conference on Computer Vision and Pattern Recognition. Piscataway: IEEE, 2017: 5967-5976.
- [14] LIU G, REDA F A, SHIH K, et al. Image in-painting for irregular holes using partial convolutions: US201916360895 [P]. 2024-03-09.
- [15] RONNEBERGER O, FISCHER P, BROX T. U-Net: convolutional networks for biomedical image segmentation [C]//International Conference on Medical Image Computing and Computer - Assisted Intervention. Berlin: Springer, 2015: 234-241.
- [16] ZENG Y H, FU J L, CHAO H Y, et al. Aggregated contextual transformations for high-resolution image inpainting[EB/OL]. (2021-04-03). <https://arxiv.org/abs/2104.01431>.
- [17] PASCAL P, REINHARD K, BRIAN A, et al. A 3D face model for pose and illumination invariant face recognition[C]//IEEE Conference on Advanced Video and Signal Based Surveillance (AVSS). Piscataway: IEEE, 2009: 296-301.
- [18] GUO Y D, ZHANG J Y, CAI J F, et al. CNN-based real-time dense face reconstruction with inverse-rendered photo-realistic face images[J]. IEEE transactions on pattern analysis and machine intelligence, 2019, 41(6): 1294-1307.
- [19] ZHU X Y, LEI Z, LIU X M, et al. Face alignment across large poses: a 3D solution[C]//IEEE Conference on Computer Vision and Pattern Recognition. Piscataway: IEEE, 2016: 146-155.
- [20] WANG Z, BOVIK A C, SHEIKH H R, et al. Image quality assessment: from error visibility to structural similarity[J]. IEEE transactions on image processing, 2004, 13(4): 600-612.

【作者简介】

程慧敏(1998—),女,福建福州人,硕士,研究方向:三维人脸重建、计算机视觉等。

姚剑敏(1978—),男,福建莆田人,博士,副研究员,研究方向:人工智能、图像处理、计算机视觉等。

(收稿日期:2023-12-21)