

Unrestricted Facial Geometry Reconstruction Using Image-to-Image Translation

Matan Sela Elad Richardson Ron Kimmel

Department of Computer Science, Technion - Israel Institute of Technology

{matansel, eladrich, ron}@cs.technion.ac.il



Figure 1: Results of the proposed method. Reconstructed geometries are shown next to the corresponding input images.

Abstract

It has been recently shown that neural networks can recover the geometric structure of a face from a single given image. A common denominator of most existing face geometry reconstruction methods is the restriction of the solution space to some low-dimensional subspace. While such a model significantly simplifies the reconstruction problem, it is inherently limited in its expressiveness. As an alternative, we propose an Image-to-Image translation network that jointly maps the input image to a depth image and a facial correspondence map. This explicit pixel-based mapping can then be utilized to provide high quality reconstructions of diverse faces under extreme expressions, using a purely geometric refinement process. In the spirit of recent approaches, the network is trained only with synthetic data, and is then evaluated on “in-the-wild” facial images. Both qualitative and quantitative analyses demonstrate the accuracy and the robustness of our approach.

1. Introduction

Recovering the geometric structure of a face is a fundamental task in computer vision with numerous applications. For example, facial characteristics of actors in realistic movies can be manually edited with facial rigs that are carefully designed for manipulating the expression [42]. While producing animation movies, tracking the geometry of an actor across multiple frames allows transferring the expression to an animated avatar [14, 8, 7]. Image-based

face recognition methods deform the recovered geometry for producing a neutralized and frontal version of the input face in a given image, reducing the variations between images of the same subject [49, 19]. As for medical applications, acquiring the structure of a face allows for fine planning of aesthetic operations and plastic surgeries, designing of personalized masks [2, 37] and even bio-printing facial organs.

Here, we focus on the recovery of the geometric structure of a face from a single facial image under a wide range of expressions and poses. This problem has been investigated for decades and most existing solutions involve one or more of the following components.

- Facial landmarks [25, 46, 32, 47] - a set of automatically detected key points on the face such as the tip of the nose and the corners of the eyes, which can guide the reconstruction process [49, 26, 1, 12, 29].
- A reference facial model - an average neutral face that is used as an initialization of optical flow or shape from shading procedures [19, 26].
- A three-dimensional morphable model - a prior low-dimensional linear subspace of plausible facial geometries which allows an efficient, yet rough, recovery of a facial structure [4, 6, 49, 36, 23, 33, 43],

While using these components can simplify the reconstruction problem, they introduce some inherent limitations. Methods that rely only on landmarks are limited to a sparse set of constrained points. Classical techniques that use a

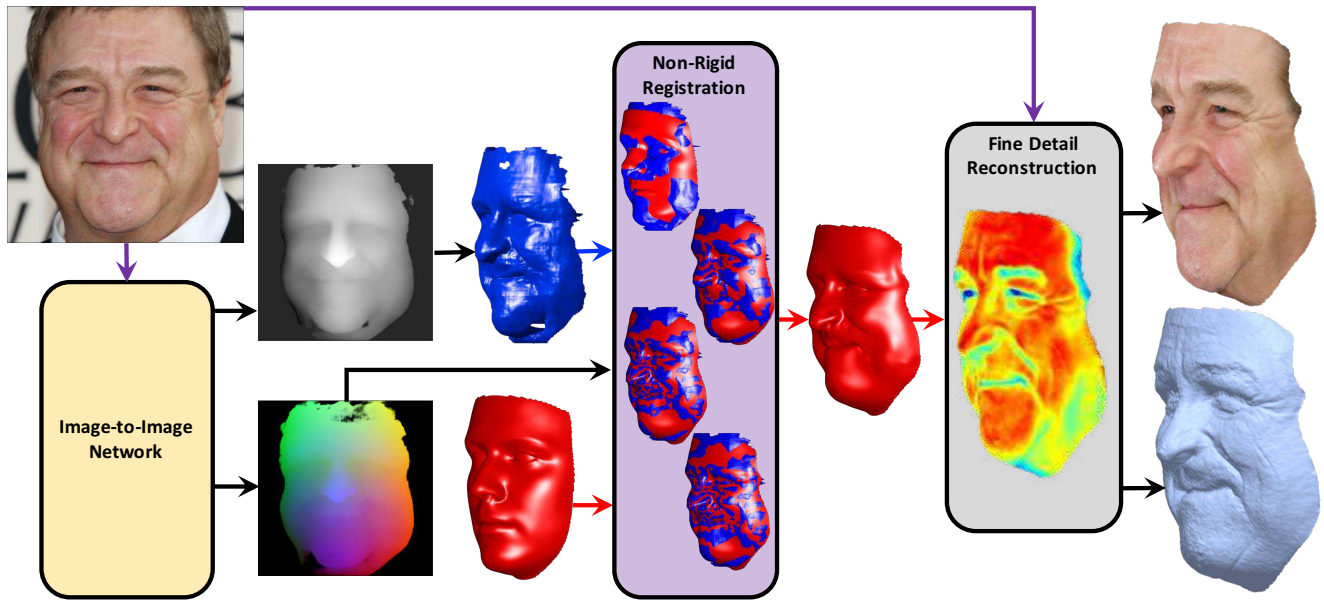


Figure 2: The algorithmic reconstruction pipeline.

reference facial model might fail to recover extreme expressions and non-frontal poses, as optical flows restrict the deformation to the image plane. The morphable model, while providing some robustness, limits the reconstruction as it can express only coarse geometries. Integrating some of these components together could mitigate the problems, yet, the underlying limitations are still manifested in the final reconstruction.

Alternatively, we propose an unrestricted approach which involves a fully convolutional network that learns to translate an input facial image to a representation containing two maps. The first map is an estimation of a depth image, while the second is an embedding of a facial template mesh in the image domain. This network is trained following the Image-to-Image translation framework of [22], where an additional normal-based loss is introduced to enhance the depth result. Similar to previous approaches, we use synthetic images for training, where the images are sampled from a wide range of facial identities, poses, expressions, lighting conditions, backgrounds and material parameters. Surprisingly, even though the network is still trained with faces that are drawn from a limited generative model, it can generalize and produce structures far and beyond the limited scope of that model. To process the raw network results, an iterative facial deformation procedure is used which combines the representations into a full facial mesh. Finally, a refinement step is applied to produce a detailed reconstruction. This novel blending of neural networks with purely geometric techniques allows us to reconstruct high-quality meshes with wrinkles and details at a mesoscopic-level from only a single image.

While using a neural network for face reconstruction was proposed in the past [33, 34, 43, 48, 24], previous methods were still limited by the expressiveness of the linear model. In [34], a second network was proposed to refine the coarse facial reconstruction, yet, it could not compensate for large geometric variations beyond the given subspace. For example, the structure of the nose was still limited by the span of a facial morphable model. By learning the unconstrained geometry directly in the image domain, we overcome this limitation, as demonstrated by both quantitative and qualitative experimental results. To further analyze the potential of the proposed representation we devise an application for translating images from one domain to another. As a case study, we transform synthetic facial images into realistic ones, enforcing our network as a loss function to preserve the geometry throughout the cross domain mapping.

The main contributions of this paper are:

- A novel formulation for predicting a geometric representation of a face from a single image, which is not restricted to a linear model.
- A purely geometric deformation and refinement procedure that utilizes the network representation to produce high quality facial reconstructions.
- A novel application of the proposed network which allows translating synthetic facial images into realistic ones, while keeping the geometric structure intact.

2. Overview

The algorithmic pipeline is presented in Figure 2. The input of the network is a facial image, and the network produces two outputs: The first is an estimated depth map aligned with the input image. The second output is a dense map from each pixel to a corresponding vertex on a reference facial mesh. To bring the results into full vertex correspondence and complete occluded parts of the face, we warp a template mesh in the three-dimensional space by an iterative non-rigid deformation procedure. Finally, a fine detail reconstruction algorithm guided by the input image recovers the subtle geometric structure of the face. Code for evaluation is available at <https://github.com/matansel/pix2vertex>.

3. Learning the Geometric Representation

There are several design choices to consider when working with neural networks. First and foremost is the training data, including the input channels, their labels, and how to gather the samples. Second is the choice of the architecture. A common approach is to start from an existing architecture [27, 39, 40, 20] and to adapt it to the problem at hand. Finally, there is the choice of the training process, including the loss criteria and the optimization technique. Next, we describe our choices for each of these elements.

3.1. The Data and its Representation

The purpose of the suggested network is to regress a geometric representation from a given facial image. This representation is composed of the following two components:

Depth Image A depth profile of the facial geometry. Indeed, for many facial reconstruction tasks providing only the depth profile is sufficient [18, 26].

Correspondence Map An embedding which allows mapping image pixels to points on a template facial model, given as a triangulated mesh. To compute this signature for any facial geometry, we paint each vertex with the x , y , and z coordinates of the corresponding point on a normalized canonical face. Then, we paint each pixel in the map

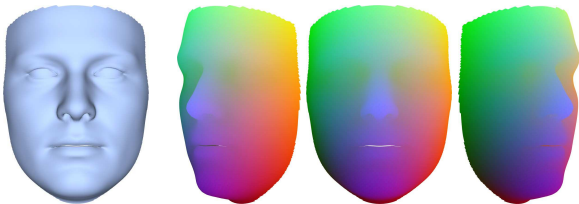


Figure 3: A reference template face presented alongside the dense correspondence signature from different viewpoints.

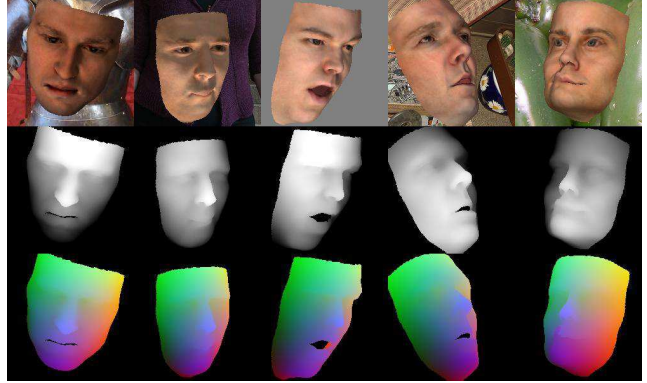


Figure 4: Training data samples alongside their representations.

with the color value of the corresponding projected vertex, see Figure 3. This feature map is a deformation agnostic representation, which is useful for applications such as facial motion capture [44], face normalization [49] and texture mapping [50]. While a similar representation was used in [34, 48] as feedback channel for an iterative network, the facial recovery was still restricted to the span of a facial morphable model.

For training the network, we adopt the same synthetic data generation procedure proposed in [33]. Each random face is generated by drawing random mesh coordinates S and texture T from a facial morphable model [4]. In practice, we draw a pair of Gaussian random vectors, α_g and α_t , and recover the synthetic face as follows

$$S = \mu_g + A_g \alpha_g$$

$$T = \mu_t + A_t \alpha_t.$$

where μ_g and μ_t are the stacked average facial geometry and texture of the model, respectively. A_g and A_t are matrices whose columns are the bases of low-dimensional linear subspaces spanning plausible facial geometries and textures, respectively. Notice that geometry basis A_g is composed to both identity and expression basis elements, as proposed in [10]. Next, we render the random textured meshes under various illumination conditions and poses, generating a dataset of synthetic facial images. As the ground-truth geometry is known for each synthetic image, one readily has the matching depth and correspondence maps to use as labels. Some examples of input images alongside their desired outputs are shown in Figure 4.

Working with synthetic data can still present some gaps when generalizing to “in-the-wild” images [9, 33], however it provides much-needed flexibility in the generation process and ensures a deterministic connection from an image to its label. Alternatively, other methods [16, 43] proposed to generate training data by employing existing reconstruction algorithms and regarding their results as ground-truth

labels. For example, Güler *et al.* [16], used a framework similar to that of [48] to match dense correspondence maps to a dataset of facial images, starting from only a sparse set of landmarks. These correspondence maps were then used as training labels for their method. Notice that such data can also be used for training our network without requiring any other modification.

3.2. Image to Geometry Translation

Pixel-wise prediction requires a proper network architecture [30, 17]. The proposed structure is inspired by the recent Image-to-Image translation framework proposed in [22], where a network was trained to map the input image to output images of various types. The architecture used there is based on the U-net [35] layout, where skip connections are used between corresponding layers in the encoder and the decoder. Additional considerations as to the network implementation are given in the supplementary.

While in [22] a combination of L_1 and adversarial loss functions were used, in the proposed framework, we chose to omit the adversarial loss. That is because unlike the problems explored in [22], our setup includes less ambiguity in the mapping. Hence, a distributional loss function is less effective, and mainly introduces artifacts. Still, since the basic L_1 loss function favors sparse errors in the depth prediction and does not account for differences between pixel neighborhoods, it is insufficient for producing fine geometric structures, see Figure 5b. Hence, we propose to augment the loss function with an additional L_1 term, which penalizes the discrepancy between the normals of the reconstructed depth and ground truth.

$$L_N(\hat{z}, z) = \|\vec{n}(\hat{z}) - \vec{n}(z)\|_1, \quad (1)$$

where $\hat{z}$ is the recovered depth, and z denotes the ground-truth depth image. During training we set $\lambda_{L_1} = 100$ and $\lambda_N = 10$, where λ_{L_1} and λ_N are the matching loss weights. Note that for the correspondence image only the L_1 loss was applied. Figure 5 demonstrates the contribution of the L_N to the quality of the depth reconstruction provided by the network.

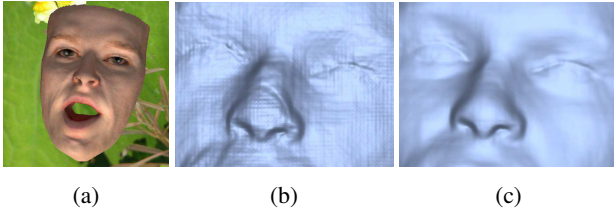


Figure 5: (a) the input image, (b) the result with only the L_1 loss function and (c) the result with the additional normals loss function. Note the artifacts in (b).

4. From Representations to a Mesh

Based on the resulting depth and correspondence we introduce an approach to translate the 2.5D representation to a 3D facial mesh. The procedure is composed of an iterative elastic deformation algorithm (4.1) followed by a fine detail recovery step driven by the input image (4.2). The resulting output is an accurate reconstructed facial mesh with a full vertex correspondence to a template mesh with fixed triangulation. This type of data is helpful for various dynamic facial processing applications, such as facial rigs, which allows creating and editing photo-realistic animations of actors. As a byproduct, this process also corrects the prediction of the network by completing domains in the face which are mistakenly classified as part of the background.

4.1. Non-Rigid Registration

Next, we describe the iterative deformation-based registration pipeline. First, we turn the depth map from the network into a mesh, by connecting neighboring pixels. Based on the correspondence map from the network, we compute the affine transformation from a template face to the mesh. This operation is done by minimizing the squared Euclidean distances between corresponding vertex pairs. Next, similar to [28], an iterative non-rigid registration process deforms the transformed template, aligning it with the mesh. Note that throughout the registration, only the template is warped, while the target mesh remains fixed. Each iteration involves the following four steps.

1. Each vertex in the template mesh, $v_i \in \mathcal{V}$, is associated with a vertex, c_i , on the target mesh, by evaluating the nearest neighbor in the correspondence embedding space. This step is different from the method described in [28], which computes the nearest neighbor in the Euclidean space. As a result, the proposed step allows registering a single template face to different facial identities with arbitrary expressions.
2. Pairs, (v_i, c_i) , which are physically distant and those whose normal directions disagree are detected and ignored in the next step.
3. The template mesh is deformed by minimizing the following energy

$$\begin{aligned} E(V, C) = & \alpha_{p2point} \sum_{(v_i, c_i) \in \mathcal{J}} \|v_i - c_i\|_2^2 \\ & + \alpha_{p2plane} \sum_{(v_i, c_i) \in \mathcal{J}} |\vec{n}(c_i)(v_i - c_i)|^2 \\ & + \alpha_{memb} \sum_{i \in \mathcal{V}} \sum_{v_j \in \mathcal{N}(v_i)} w_{i,j} \|v_i - v_j\|_2^2, \end{aligned} \quad (2)$$

where, $w_{i,j}$ is the weight corresponding to the biharmonic Laplacian operator (see [21, 5]), $\vec{n}(c_i)$ is the

normal of the corresponding vertex at the target mesh c_i , $\mathcal{J}$ is the set of the remaining associated vertex pairs (v_i, c_i) , and $\mathcal{N}(v_i)$ is the set 1-ring neighboring vertices about the vertex v_i . Notice that the first term above is the sum of squared Euclidean distances between matches. The second term is the distance from the point v_i to the tangent plane at the corresponding point of the target mesh. The third term quantifies the stiffness of the mesh.

4. If the motion of the template mesh between the current iteration and the previous one is below a fixed threshold, we divide the weight α_{memb} by two. This relaxes the stiffness term and allows a greater deformation in the next iteration.

This iterative process terminates when the stiffness weight is below a given threshold. Further implementation information and parameters of the registration process are provided in the supplementary material. The resulting output of this phase is a deformed template with fixed triangulation, which contains the overall facial structure recovered by the network, yet, is smoother and complete, see the third column of Figure 9.

4.2. Fine Detail Reconstruction

Although the network already recovers some fine geometric details, such as wrinkles and moles, across parts of the face, a geometric approach can reconstruct details at a finer level, on the entire face, independently of the resolution. Here, we propose an approach motivated by the passive-stereo facial reconstruction method suggested in [3]. The underlying assumption here is that subtle geometric structures can be explained by local variations in the image domain. For some skin tissues, such as nevi, this assumption is inaccurate as the intensity variation results from the albedo. In such cases, the geometric structure would be wrongly modified. Still, for most parts of the face, the reconstructed details are consistent with the actual variations in depth.

The method begins from an interpolated version of the deformed template. Each vertex $v \in \mathcal{V}_D$ is painted with the intensity value of the nearest pixel in the image plane. Since we are interested in recovering small details, only the high spatial frequencies, $\mu(v)$, of the texture, $\tau(v)$, are taken into consideration in this phase. For computing this frequency band, we subtract the synthesized low frequencies from the original intensity values. This low-pass filtered part can be computed by convolving the texture with a spatially varying Gaussian kernel in the image domain, as originally proposed. In contrast, since this convolution is equivalent to computing the heat distribution upon the shape after time dt , where the initial heat profile is the original texture, we

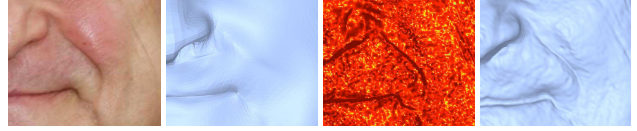


Figure 6: Mesoscopic displacement. From left to right: an input image, the shape after the iterative registration, the high-frequency part of the texture - $\mu(v)$, and the final shape.

propose to compute $\mu(v)$ as

$$\mu(v) = \tau(v) - (I - dt \cdot \Delta_g)^{-1} \tau(v), \quad (3)$$

where I is the identity matrix, Δ_g is the cotangent weight discrete Laplacian operator for triangulated meshes [31], and dt is a scalar proportional to the cut-off frequency of the filter.

Next, we displace each vertex along its normal direction such that $v' = v + \delta(v)\vec{n}(v)$. The step size of the displacement, $\delta(v)$, is a combination of a data-driven term, $\delta_\mu(v)$, and a regularization one, $\delta_s(v)$. The data-driven term is guided by the high-pass filtered part of the texture, $\mu(v)$. In practice, we require the local differences in the geometry to be proportional to the local variation in the high frequency band of the texture. For each vertex v , with a normal $\vec{n}(v)$, and a neighboring vertex v_i , the data-driven term is given by

$$\delta_\mu(v) = \frac{\sum_{v_i \in \mathcal{N}(v)} \alpha_{(v,v_i)} (\mu(v) - \mu(v_i)) \left(1 - \frac{|\langle v - v_i, \vec{n}(v) \rangle|}{\|v - v_i\|}\right)}{\sum_{v_i \in \mathcal{N}(v)} \alpha_{(v,v_i)}}, \quad (4)$$

where $\alpha_{(v,v_i)} = \exp(-\|v - v_i\|)$. For further explanation of Equation 4, we refer the reader to the supplementary material of this paper or the implementation details of [3].

Since we move each vertex along the normal direction, triangles could intersect each other, particularly in domains of high curvature. To reduce the probability of such collisions, a regularizing displacement field, $\delta_s(v)$, is added. This term is proportional to the mean curvature of the original surface, and is equivalent to a single explicit mesh fairing step [11]. The final surface modification is given by

$$v' = v + (\eta \delta_\mu(v) + (1 - \eta) \delta_s(v)) \cdot \vec{n}(v), \quad (5)$$

for some constant $\eta \in [0, 1]$. A demonstration of the results before and after this step is presented in Figure 6

5. Experiments

Next, we present evaluations on both the proposed network and the pipeline as a whole, and comparison to different prominent methods of single image based facial reconstruction [26, 49, 34].

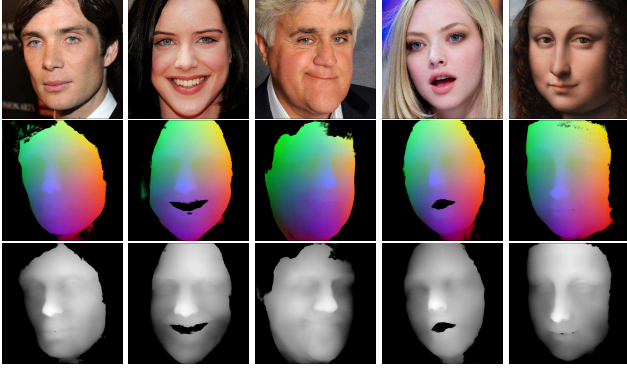


Figure 7: Network Output.

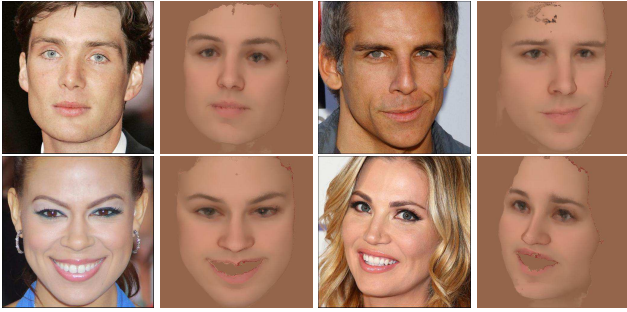


Figure 8: Texture mapping via the embedding.

5.1. Qualitative Evaluation

The first component of our algorithm is an Image-to-Image network. In Figure 7, we show samples of output maps produced by the proposed network. Although the network was trained with synthetic data, with simple random backgrounds (see Figure 4), it successfully separates the hair and background from the face itself and learns the corresponding representations. To qualitatively assess the accuracy of the correspondence, we present a visualization where an average facial texture is mapped to the image plane via the predicted embedding, see Figure 8, this shows how the network successfully learns to represent the facial structure. Next, in Figure 9 we show the reconstruction of the network, alongside the registered template and the final shape. Notice how the structural information retrieved by the network is preserved through the geometric stages. Figure 10 shows a qualitative comparison between the proposed method and others. One can see that our method better matches the global structure, as well as the facial details. To better perceive these differences, see Figure 11. Finally, to demonstrate the limited expressiveness of the 3DMM space compared to our method, Figure 12 presents our registered template next to its projection onto the 3DMM space. This clearly shows that our network is able to learn structures which are not spanned by the 3DMM model.

5.2. Quantitative Evaluation

For a quantitative comparison, we used the first 200 subjects from the BU-3DFE dataset [45], which contains facial images aligned with ground truth depth images. Each method provides its own estimation for the depth image alongside a binary mask, representing the valid pixels to be taken into account in the evaluation. Obviously, since the problem of reconstructing depth from a single image is ill-posed, the estimation needs to be judged up to global scaling and transition along the depth axis. Thus, we compute these parameters using the Random Sample Consensus (RANSAC) approach [13], for normalizing the estimation according to the ground truth depth. This significantly reduces the absolute error of each method as the global parameter estimation is robust to outliers. Note that the parameters of the RANSAC were identical for all the methods and samples. The results of this comparison are given in Table 1, where the units are given in terms of the percentile of the ground-truth depth range. As a further analysis of the reconstruction accuracy, we computed the mean absolute error of each method based on expressions, see Table 2.

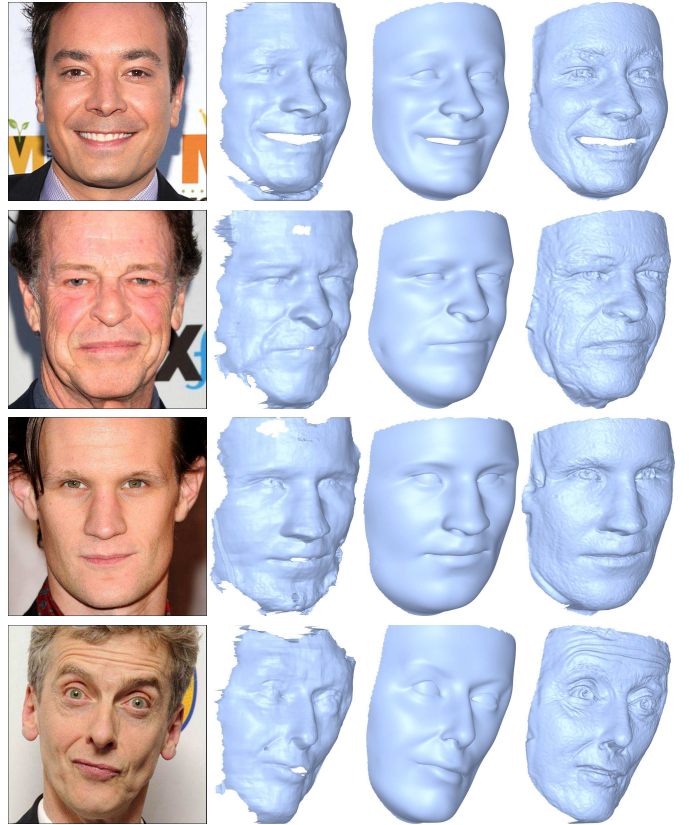


Figure 9: The reconstruction stages. From left to right: the input image, the reconstruction of the network, the registered template and the final shape.

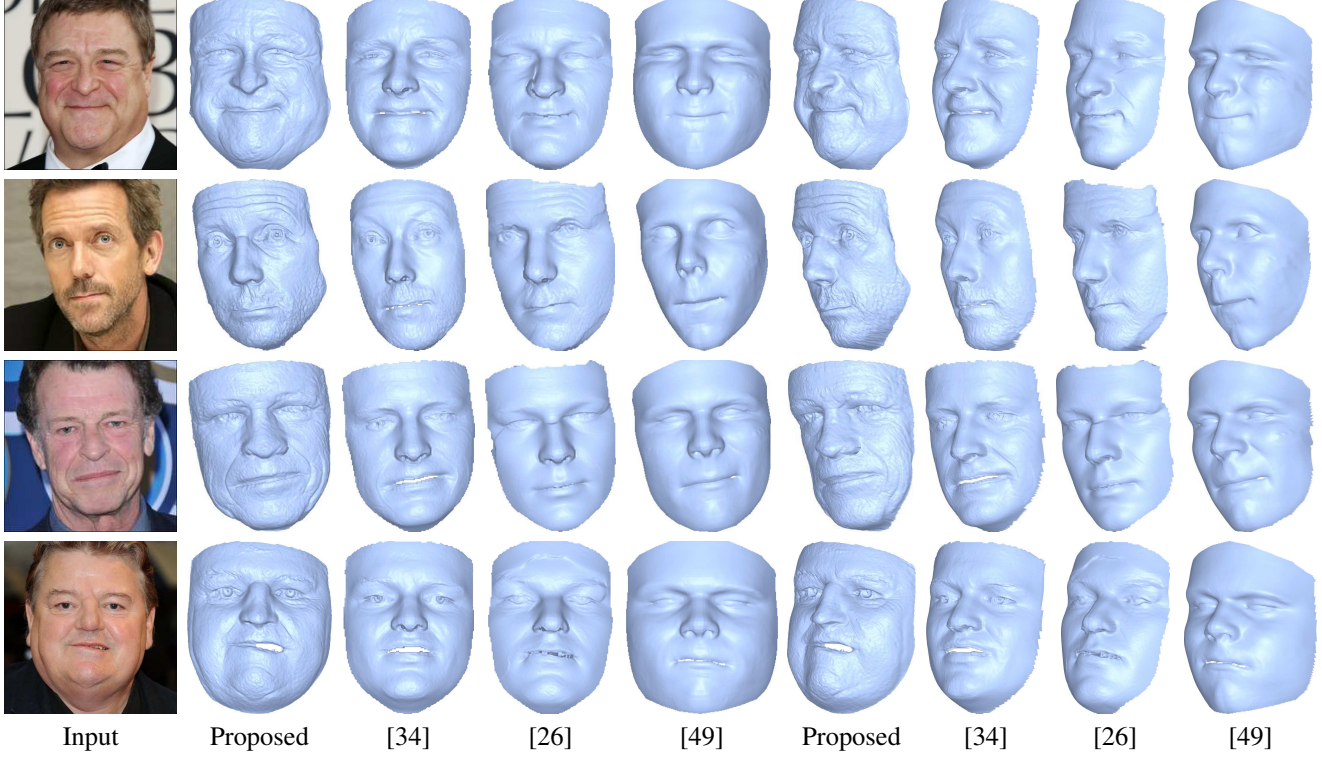


Figure 10: Qualitative comparison. Input images are presented alongside the reconstructions of the different methods.

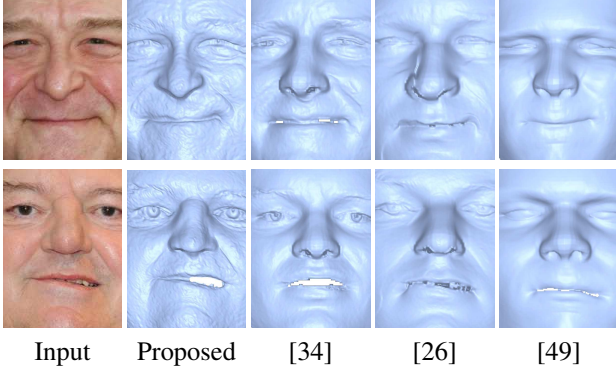


Figure 11: Zoomed qualitative result of first and fourth subjects from Figure 10.

	Mean Err.	Std Err.	Median Err.	90% Err.
[26]	3.89	4.14	2.94	7.34
[49]	3.85	3.23	2.93	7.91
[34]	3.61	2.99	2.72	6.82
Ours	3.51	2.69	2.65	6.59

Table 1: Quantitative evaluation on the BU-3DFE Dataset. From left to right: the absolute depth errors evaluated by mean, standard deviation, median and the average ninety percent largest error.

	AN	DI	FE	HA	NE	SA	SU
[26]	3.47	4.03	3.94	4.30	3.43	3.52	4.19
[49]	4.00	3.93	3.91	3.70	3.76	3.61	3.96
[34]	3.42	3.46	3.64	3.41	4.22	3.59	4.00
Ours	3.67	3.34	3.36	3.01	3.17	3.37	4.41

Table 2: The mean error by expression. From left to right: Anger, Disgust, Fear, Happy, Neutral, Sad, Surprise.

5.3. The Network as a Geometric Constraint

As demonstrated by the results, the proposed network successfully learns both the depth and the embedding representations for a variety of images. This representation is the key part behind the reconstruction pipeline. However, it can also be helpful for other face-related tasks. As an example, we show that the network can be used as a geometric constraint for facial image manipulations, such as transforming synthetic images into realistic ones. This idea

is based on recent advances in applying Generative Adversarial Networks (GAN) [15] for domain adaption tasks [41].

In the basic GAN framework, a Generator Network (G) learns to map from the source domain, $\mathcal{D}_S$, to the target domain $\mathcal{D}_T$, where a Discriminator Network (D) tries to dis-

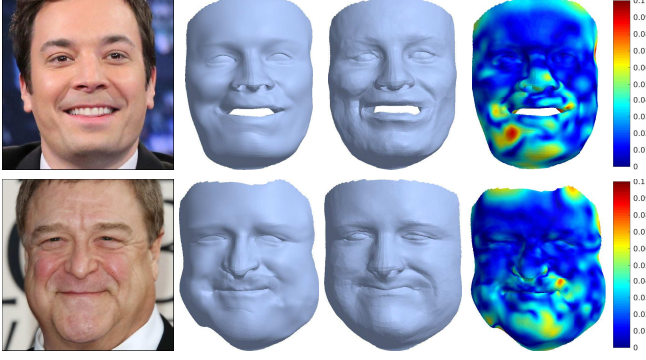


Figure 12: 3DMM Projection. From left to right: the input image, the registered template, the projected mesh and the projection error.

tinguish between generated images and samples from the target domain, by optimizing the following objective

$$\min_G \max_D V(D, G) = \mathbb{E}_{y \sim \mathcal{D}_T} [\log D(y)] + \mathbb{E}_{x \sim \mathcal{D}_S} [\log (1 - D(G(x)))] \quad (6)$$

Theoretically, this framework could also translate images from the synthetic domain into the realistic one. However, it does not guarantee that the underlying geometry of the synthetic data is preserved throughout that transformation. That is, the generated image might look realistic, but have a completely different facial structure from the synthetic input. To solve that potential inconsistency, we suggest to involve the proposed network as an additional loss function on the output of the generator.

$$L_{Geom}(x) = \|Net(x) - Net(G(x))\|_1, \quad (7)$$

where $Net(\cdot)$ represents the operation of the introduced network. Note that this is feasible, thanks to the fact that the proposed network is fully differentiable. The additional geometric fidelity term forces the generator to learn a mapping that makes a synthetic image more realistic while keeping the underlying geometry intact. This translation process could potentially be useful for data generation procedures, similarly to [38]. Some successful translations are visualized in Figure 13. Notice that the network implicitly learns to add facial hair and teeth, and modify the texture and shading, without changing the facial structure. As demonstrated by this analysis, the proposed network learns a strong representation that has merit not only for reconstruction, but for other tasks as well.

6. Limitations

One of the core ideas of this work was a model-free approach, where the solution space is not restricted by a low dimensional subspace. Instead, the Image-to-Image

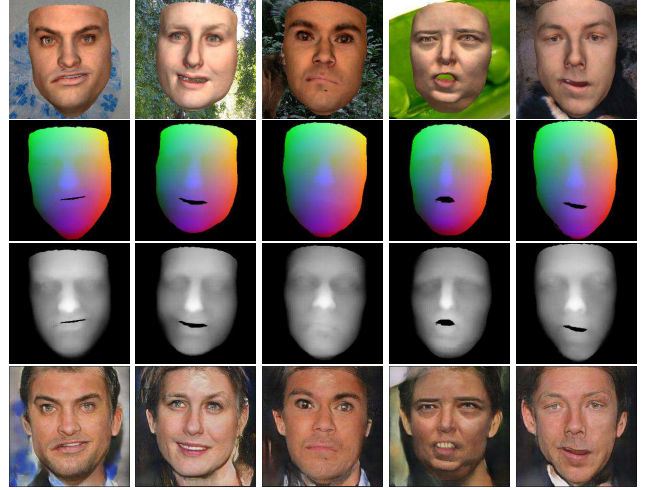


Figure 13: Translation results. From top to bottom: synthetic input images, the correspondence and the depth maps recovered by the network, and the transformed result.

network represents the solution in the extremely high-dimensional image domain. This structure is learned from synthetic examples, and shown to successfully generalize to “in-the-wild” images. Still, facial images that significantly deviate from our training domain are challenging, resulting in missing areas and errors inside the representation maps. More specifically, our network has difficulty handling extreme occlusions such as sunglasses, hands or beards, as these were not seen in the training data. Similarly to other methods, reconstructions under strong rotations are also not well handled. Reconstructions under such scenarios are shown in the supplementary material. Another limiting factor of our pipeline is speed. While the suggested network by itself can be applied efficiently, our template registration step is currently not optimized for speed and can take a few minutes to converge.

7. Conclusion

We presented an unrestricted approach for recovering the geometric structure of a face from a single image. Our algorithm employs an Image-to-Image network which maps the input image to a pixel-based geometric representation, followed by geometric deformation and refinement steps. The network is trained only by synthetic facial images, yet, is capable of reconstructing real faces. Using the network as a loss function, we propose a framework for translating synthetic facial images into realistic ones while preserving the geometric structure.

Acknowledgments

We would like to thank Roy Or-EI for the helpful discussions and comments.

References

- [1] O. Aldrian and W. A. Smith. A linear approach of 3D face shape and texture recovery using a 3d morphable model. In *Proceedings of the British Machine Vision Conference*, pages 75–1, 2010.
- [2] I. Amirav, A. S. Luder, A. Halamish, D. Raviv, R. Kimmel, D. Waisman, and M. T. Newhouse. Design of aerosol face masks for children using computerized 3d face analysis. *Journal of aerosol medicine and pulmonary drug delivery*, 27(4):272–278, 2014.
- [3] T. Beeler, B. Bickel, P. Beardsley, B. Sumner, and M. Gross. High-quality single-shot capture of facial geometry. In *ACM SIGGRAPH 2010 Papers*, SIGGRAPH '10, pages 40:1–40:9, New York, NY, USA, 2010. ACM.
- [4] V. Blanz and T. Vetter. A morphable model for the synthesis of 3D faces. In *Proceedings of the 26th annual conference on Computer graphics and interactive techniques*, pages 187–194. ACM Press/Addison-Wesley Publishing Co., 1999.
- [5] M. Botsch and O. Sorkine. On linear variational surface deformation methods. *IEEE Transactions on Visualization and Computer Graphics*, 14(1):213–230, Jan 2008.
- [6] P. Breuer, K.-I. Kim, W. Kienzle, B. Scholkopf, and V. Blanz. Automatic 3D face reconstruction from single images or video. In *Automatic Face & Gesture Recognition, 2008. FG'08. 8th IEEE International Conference on*, pages 1–8. IEEE, 2008.
- [7] C. Cao, D. Bradley, K. Zhou, and T. Beeler. Real-time high-fidelity facial performance capture. *ACM Transactions on Graphics (TOG)*, 34(4):46, 2015.
- [8] C. Cao, Y. Weng, S. Lin, and K. Zhou. 3D shape regression for real-time facial animation. *ACM Transactions on Graphics (TOG)*, 32(4):41, 2013.
- [9] W. Chen, H. Wang, Y. Li, H. Su, D. Lischinsk, D. Cohen-Or, B. Chen, et al. Synthesizing training images for boosting human 3D pose estimation. *arXiv preprint arXiv:1604.02703*, 2016.
- [10] B. Chu, S. Romdhani, and L. Chen. 3d-aided face recognition robust to expression and pose variations. In *2014 IEEE Conference on Computer Vision and Pattern Recognition*, pages 1907–1914. IEEE, 2014.
- [11] M. Desbrun, M. Meyer, P. Schröder, and A. H. Barr. Implicit fairing of irregular meshes using diffusion and curvature flow. In *Proceedings of the 26th Annual Conference on Computer Graphics and Interactive Techniques*, SIGGRAPH '99, pages 317–324, New York, NY, USA, 1999. ACM Press/Addison-Wesley Publishing Co.
- [12] P. Dou, Y. Wu, S. K. Shah, and I. A. Kakadiaris. Robust 3D face shape reconstruction from single images via two-fold coupled structure learning. In *Proc. British Machine Vision Conference*, pages 1–13, 2014.
- [13] M. A. Fischler and R. C. Bolles. Random sample consensus: a paradigm for model fitting with applications to image analysis and automated cartography. *Communications of the ACM*, 24(6):381–395, 1981.
- [14] P. Garrido, M. Zollhöfer, D. Casas, L. Valgaerts, K. Varanasi, P. Pérez, and C. Theobalt. Reconstruction of personalized 3D face rigs from monocular video. *ACM Transactions on Graphics (TOG)*, 35(3):28, 2016.
- [15] I. Goodfellow, J. Pouget-Abadie, M. Mirza, B. Xu, D. Warde-Farley, S. Ozair, A. Courville, and Y. Bengio. Generative adversarial nets. In *Advances in Neural Information Processing Systems*, pages 2672–2680, 2014.
- [16] R. A. Güler, G. Trigeorgis, E. Antonakos, P. Snape, S. Zafeiriou, and I. Kokkinos. Densereg: Fully convolutional dense shape regression in-the-wild. *arXiv preprint arXiv:1612.01202*, 2016.
- [17] B. Hariharan, P. Arbeláez, R. Girshick, and J. Malik. Hypercolumns for object segmentation and fine-grained localization. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pages 447–456, 2015.
- [18] T. Hassner. Viewing real-world faces in 3d. In *Proceedings of the IEEE International Conference on Computer Vision*, pages 3607–3614, 2013.
- [19] T. Hassner, S. Harel, E. Paz, and R. Enbar. Effective face frontalization in unconstrained images. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pages 4295–4304, 2015.
- [20] K. He, X. Zhang, S. Ren, and J. Sun. Deep residual learning for image recognition. In *The IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, June 2016.
- [21] B. T. Helenbrook. Mesh deformation using the biharmonic operator. *International journal for numerical methods in engineering*, 56(7):1007–1021, 2003.
- [22] P. Isola, J.-Y. Zhu, T. Zhou, and A. A. Efros. Image-to-image translation with conditional adversarial networks. *arXiv preprint arXiv:1611.07004*, 2016.
- [23] L. Jiang, J. Zhang, B. Deng, H. Li, and L. Liu. 3d face reconstruction with geometry details from a single image. *arXiv preprint arXiv:1702.05619*, 2017.
- [24] A. Jourabloo and X. Liu. Large-pose face alignment via cnn-based dense 3D model fitting. In *The IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, June 2016.
- [25] V. Kazemi and J. Sullivan. One millisecond face alignment with an ensemble of regression trees. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pages 1867–1874, 2014.
- [26] I. Kemelmacher-Shlizerman and R. Basri. 3D face reconstruction from a single image using a single reference face shape. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 33(2):394–405, 2011.
- [27] A. Krizhevsky, I. Sutskever, and G. E. Hinton. Imagenet classification with deep convolutional neural networks. In *Advances in neural information processing systems*, pages 1097–1105, 2012.
- [28] H. Li. *Animation Reconstruction of Deformable Surfaces*. PhD thesis, ETH Zurich, November 2010.
- [29] F. Liu, D. Zeng, J. Li, and Q. Zhao. Cascaded regressor based 3D face reconstruction from a single arbitrary view image. *arXiv preprint arXiv:1509.06161*, 2015.
- [30] J. Long, E. Shelhamer, and T. Darrell. Fully convolutional networks for semantic segmentation. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pages 3431–3440, 2015.

- [31] M. Meyer, M. Desbrun, P. Schröder, A. H. Barr, et al. Discrete differential-geometry operators for triangulated 2-manifolds. *Visualization and mathematics*, 3(2):52–58, 2002.
- [32] S. Ren, X. Cao, Y. Wei, and J. Sun. Face alignment at 3000 fps via regressing local binary features. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pages 1685–1692, 2014.
- [33] E. Richardson, M. Sela, and R. Kimmel. 3D face reconstruction by learning from synthetic data. In *3D Vision (3DV), 2016 International Conference on*, pages 460–469. IEEE, 2016.
- [34] E. Richardson, M. Sela, R. Or-El, and R. Kimmel. Learning detailed face reconstruction from a single image. *arXiv preprint arXiv:1611.05053*, 2016.
- [35] O. Ronneberger, P. Fischer, and T. Brox. U-net: Convolutional networks for biomedical image segmentation. In *International Conference on Medical Image Computing and Computer-Assisted Intervention*, pages 234–241. Springer, 2015.
- [36] S. Saito, L. Wei, L. Hu, K. Nagano, and H. Li. Photorealistic facial texture inference using deep neural networks. *arXiv preprint arXiv:1612.00523*, 2016.
- [37] M. Sela, N. Toledo, Y. Honen, and R. Kimmel. Customized facial constant positive air pressure (cpap) masks. *arXiv preprint arXiv:1609.07049*, 2016.
- [38] A. Shrivastava, T. Pfister, O. Tuzel, J. Susskind, W. Wang, and R. Webb. Learning from simulated and unsupervised images through adversarial training. *arXiv preprint arXiv:1612.07828*, 2016.
- [39] K. Simonyan and A. Zisserman. Very deep convolutional networks for large-scale image recognition. *arXiv preprint arXiv:1409.1556*, 2014.
- [40] C. Szegedy, W. Liu, Y. Jia, P. Sermanet, S. Reed, D. Anguelov, D. Erhan, V. Vanhoucke, and A. Rabinovich. Going deeper with convolutions. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pages 1–9, 2015.
- [41] Y. Taigman, A. Polyak, and L. Wolf. Unsupervised cross-domain image generation. *arXiv preprint arXiv:1611.02200*, 2016.
- [42] J. Thies, M. Zollhofer, M. Stamminger, C. Theobalt, and M. Nießner. Face2face: Real-time face capture and reenactment of rgb videos. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pages 2387–2395, 2016.
- [43] A. T. Tran, T. Hassner, I. Masi, and G. Medioni. Regressing robust and discriminative 3d morphable models with a very deep neural network. *arXiv preprint arXiv:1612.04904*, 2016.
- [44] T. Weise, S. Bouaziz, H. Li, and M. Pauly. Realtime performance-based facial animation. In *ACM Transactions on Graphics (TOG)*, volume 30, page 77. ACM, 2011.
- [45] L. Yin, X. Wei, Y. Sun, J. Wang, and M. J. Rosato. A 3d facial expression database for facial behavior research. In *Automatic face and gesture recognition, 2006. FGR 2006. 7th international conference on*, pages 211–216. IEEE, 2006.
- [46] Z. Zhang, P. Luo, C. C. Loy, and X. Tang. Facial landmark detection by deep multi-task learning. In *European Conference on Computer Vision*, pages 94–108. Springer, 2014.
- [47] E. Zhou, H. Fan, Z. Cao, Y. Jiang, and Q. Yin. Extensive facial landmark localization with coarse-to-fine convolutional network cascade. In *Proceedings of the IEEE International Conference on Computer Vision Workshops*, pages 386–391, 2013.
- [48] X. Zhu, Z. Lei, X. Liu, H. Shi, and S. Z. Li. Face alignment across large poses: A 3d solution. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pages 146–155, 2016.
- [49] X. Zhu, Z. Lei, J. Yan, D. Yi, and S. Z. Li. High-fidelity pose and expression normalization for face recognition in the wild. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pages 787–796, 2015.
- [50] G. Zigelman, R. Kimmel, and N. Kiryati. Texture mapping using surface flattening via multidimensional scaling. *IEEE Transactions on Visualization and Computer Graphics*, 8(2):198–207, 2002.