

ARAPReg: An As-Rigid-As Possible Regularization Loss for Learning Deformable Shape Generators

Qixing Huang*
UT Austin
huangqx@cs.utexas.edu

Xiangru Huang*
UT Austin & MIT
xiangruhuang816@gmail.com

Bo Sun*
UT Austin
bosun@cs.utexas.edu

Zaiwei Zhang
UT Austin
zaiweizhang@utexas.edu

Junfeng Jiang
Hohai University
jiangjf.hhu@gmail.com

Chandrajit Bajaj
UT Austin
bajaj@cs.utexas.edu

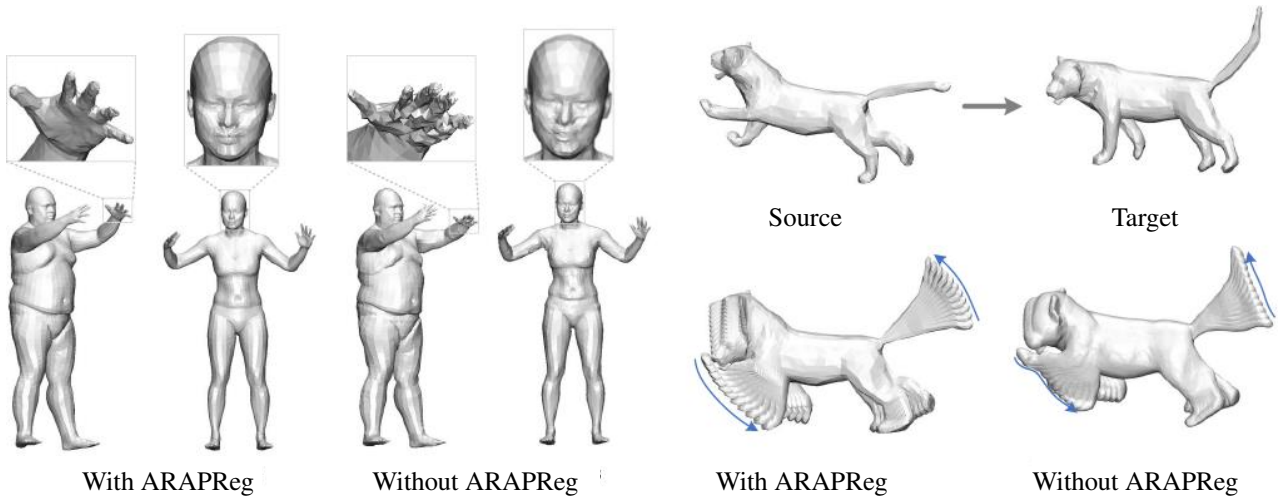


Figure 1: Our approach learns a shape generator from a collection of deformable shapes. The shape generator is trained with a novel as-rigid-as-possible regularization (ARAPReg) loss that promotes the preservation of multi-scale shape features.

Abstract

This paper introduces an unsupervised loss for training parametric deformation shape generators. The key idea is to enforce the preservation of local rigidity among the generated shapes. Our approach builds on an approximation of the as-rigid-as possible (or ARAP) deformation energy. We show how to develop the unsupervised loss via a spectral decomposition of the Hessian of the ARAP energy. Our loss nicely decouples pose and shape variations through a robust norm. The loss admits simple closed-form expressions. It is easy to train and can be plugged into any standard generation models, e.g., variational auto-encoder (VAE) and auto-decoder (AD). Experimental results show that our approach outperforms existing shape generation approaches considerably on public benchmark datasets of various shape categories such as human, animal and bone. Our code and data are available at <https://github.com/GitBoSun/ARAPReg>.

1. Introduction

This paper considers learning a parametric mesh generator from a deformable shape collection with shapes that exhibit the same topology but undergo large geometric variations (see examples below of a deforming human, animal, and bone). This problem arises in numerous visual computing and relevant fields such as recovery of neural morphogenesis, data-driven shape reconstruction, and image-based reconstruction, to name just a few (c.f. [54]).

Deformable shapes differ from many other visual objects (e.g., images and videos) because there are natural constraints underlying the shape space. One such example is the local rigidity constraint; namely, corresponding surface patches among neighboring shapes in the shape space undergo approximately rigid transformations. This constraint manifests the preservation of geometric features (e.g., facial features of humans and toes of animals) among local neighborhoods of the underlying shape space. An interesting problem thus is the use of this constraint to train shape generators from a collection of training shapes, where

the local rigidity constraint accurately and efficiently propagates features of the training shapes to new synthetic shapes produced by the generator.

In this paper, we study how to model the local rigidity constraint as an unsupervised loss functional for generative modeling. The proposed loss can be combined with standard mesh generators such as variational auto-encoders (VAEs) [44, 28, 39, 7] and auto-decoders (ADs) [59, 61]. A key property of our loss functional is that it is consistent with other training losses. This property offers multiple advantages. For example, the learned generator is insensitive to the tradeoff parameters among the loss terms. As another example, the training procedure converges faster than the setting where loss terms may compete against each other.

Our approach, called ARAPReg, builds on the established as-rigid-as-possible (or ARAP) deformation model [43, 49, 55] that measures the non-rigid deformation between two shapes. Its key ingredients include use of the Hessian of the ARAP deformation model to derive an explicit regularizer for the Jacobian of the shape generator and a robust norm on the Hessian to model pose and shape variations of deformable shapes. The outcome is a simple closed-form formulation for training mesh generators. ARAPReg differs from prior works that enforce ARAP losses between synthetic shapes and a base shape [17, 27, 63], that may introduce competing losses when the underlying shape space has large deformations.

We have evaluated ARAPReg across a variety of public benchmark datasets such as DFAUST [5], SMAL [66], and an in-house benchmark dataset of Bone. The evaluations include both generator settings of VAE and AD. Experimental results show that ARAPReg leads to considerable performance gains across state-of-the-art deformable shape generators both qualitatively and quantitatively. As shown in Figure 1 for example, the interpolated shapes using ARAPReg greatly preserve the local geometric details of the generated shapes and avoids unrealistic shape poses.

2. Related Works

This section organizes the relevant works into three groups, namely, 3D generative models, regularization for generative modeling, and shape space modeling.

3D generative models. Learning 3D generative models relies on developing suitable 3D representations to encode 3D models into vectorized forms. Examples include volumetric grid [53, 45, 12, 34, 32, 19], implicit surfaces [35, 9], point clouds [1, 57, 56, 25], meshes [22, 28, 14], parametric surfaces [16, 31], spherical representations [10, 13, 8], geometric arrangements [47, 62], and multi-views [30].

This paper is mostly relevant to generative models under the mesh representation, which falls into four categories. The first category of approaches [44, 48, 28, 46, 38] is based on defining variational auto-encoders on meshes. A typical strategy is to treat triangular meshes as graphs and define convolution and deconvolution operations to synthesize triangular meshes (c.f. [48, 28, 46]). [44] introduced a geometric encoding scheme that operates in the gradient do-

main. The second category of approaches builds upon recurrent procedures for geometric synthesis. This methodology has been extensively applied for primitive-based assembly [26, 40, 41, 65]. [18] extended this approach to meshes, in which edge contraction operations are applied recursively. The third category of approaches [60, 51] deforms a base mesh to generate new meshes, where the deformation is learned from data. The last category utilizes surface parameterization [42, 31, 16, 4].

While these approaches focused on adopting generative modeling methodologies under the mesh setting, ARAPReg studies the novel problem of explicitly enforcing an ARAP loss among synthetic shapes with similar latent codes.

Regularization for generative modeling. Regularization losses have been explored in prior works for 3D generative modeling. In [36], Peebles et al. studied a Hessian regularization term for learning generative image models. A spectral regularization loss is introduced in [2] for 3D generative modeling. Several works [52, 50, 21, 3] studied geometric regularizations for image-based reconstruction. In contrast, ARAPReg focuses on regularization terms that are consistent with other terms. Several other works [17, 27, 63] employed ARAP losses between any synthetic shapes with a base shape. The novelty of ARAPReg is that it is consistent with other loss terms even when the underlying shape space presents large deformations. The reason is that the local rigidity constraint is only enforced among neighboring shapes in the underlying shape space. Our initial experiments show that enforcing ARAP losses between synthetic shapes and a base shape leads to worse results than dropping the ARAP losses.

Shape space modeling. Finally, ARAPReg is relevant to early works on modeling tangent spaces of shape manifolds [23, 20, 58]. However, unlike the applications in shape interpolation [23], shape segmentation [20], and mesh-based geometric design [37, 58], ARAPReg focuses on devising an unsupervised loss for network training.

3. Overview

Following [39, 7, 64], we are interested in learning a mesh generator that takes a latent code as input and outputs the vertex positions of a triangular mesh with given mesh connectivity (See Figure 2). Formally speaking, we denote this mesh generator as

$$g^\theta : \mathcal{Z} := \mathcal{R}^k \rightarrow \mathcal{R}^{3n}.$$

Here $\mathcal{Z}$ represents the latent space, and $\mathcal{R}^{3n}$ encodes the vector that concatenates the vertex positions, i.e., n is the number of vertices. We organize the remainder of this paper as follows.

In Section 4, we introduce the key contribution of this paper, ARAPReg, an unsupervised loss for training g^θ . The loss only requires a prior distribution of the latent space $\mathcal{Z}$. In this paper, we assume the prior distribution is the Normal distribution $\mathcal{N}_k$ of dimension k . The key idea of ARAPReg is to ensure that the local rigidity constraint is pre-

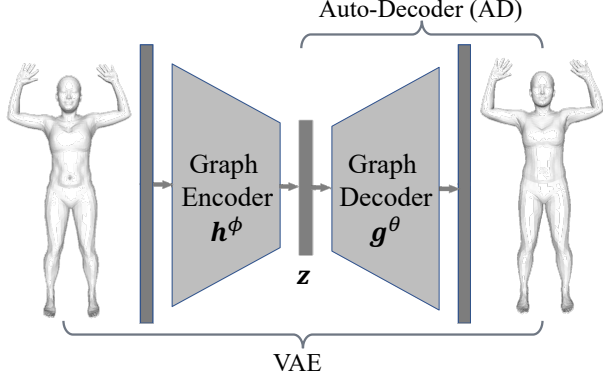


Figure 2: We consider multiple standard shape generators, including Variational Auto-Encoder (VAE) and Auto-Decoder (AD). The graph encoder h^ϕ maps the input mesh g to a latent parameter $h^\phi(z)$. The graph decoder maps a latent parameter z to the out mesh $g^\theta(z)$.

served among neighboring generated shapes in the underlying shape space. As illustrated in the left part of Figure 1, the goal of this loss is to significantly improve the generalization behavior of the mesh generator, e.g., preserving multi-scale geometric details.

In Section 5, we discuss how to plug this unsupervised loss into standard shape generation models based on VAE and AD.

4. Formulation of the ARAPReg Loss

Formulating the preservation of local rigidity is quite challenging because the resulting loss term has to be simple enough to facilitate network training. One straightforward approach is to enforce the local rigidity constraint between a generated shape $g^\theta(z)$ and its perturbation $g^\theta(z + dz)$. Here dz characterizes an infinitesimal displacement in the parameter space. However, this approach requires sampling a lot of shape pairs. Besides, typical formulations of shape deformations between $g^\theta(z)$ and $g^\theta(z + dz)$ require solving optimization problems that are computationally expensive (c.f. [6]).

ARAPReg stitches several novel ideas to derive a simple unsupervised loss that does not adversely compete with typical losses used in generative modeling (See Section 5.1).

4.1. Step I: Decoupling Smoothness and Jacobian regularization

First, ARAPReg decouples the enforcement of local rigidity into two terms. The first term enforces the smoothness of the generator. This smoothness penalty enables the second term, which formulates the preservation of local rigidity as potentials on the Jacobian of the generator, i.e.,

$$\frac{\partial g^\theta}{\partial z}(z) \in \mathcal{R}^{(3n) \times k}.$$

Specifically, we define the unsupervised loss as

$$\mathcal{L}_{reg}(\theta) := E_{z \sim \mathcal{N}_k} \left(E_{\delta z \sim s, \mathcal{N}_k} \|g^\theta(z + \delta z) - 2g^\theta(z) + g^\theta(z - \delta z)\|^2 + \lambda_R \cdot r_R(g^\theta(z), \frac{\partial g^\theta}{\partial z}(z)) \right), \quad (1)$$

where the first term promotes the smoothness of the generator g^θ ; s is a hyper-parameter of ARAPReg. Note that unlike enforcing

$$g^\theta(z + \delta z) \approx g^\theta(z) + \frac{\partial g^\theta}{\partial z}(z) \cdot \delta z, \quad (2)$$

the formulation in (1) does not involve the first-order derivatives of g . It follows that network training is more efficient as it only requires computing the first-order derivatives of g . On the other hand, it penalizes the second-order derivatives of g^θ . It therefore implicitly enforces (2). The second term $r_R(g^\theta(z), \frac{\partial g^\theta}{\partial z}(z))$ in (1), which will be defined shortly, formulates the regularization loss concerning the generated mesh $g^\theta(z)$ and infinitesimal perturbations specified by the Jacobian $\frac{\partial g^\theta}{\partial z}(z)$ (See Figure 3). λ_R is another hyper-parameter of ARAPReg.

In other words, instead of enforcing the local rigidity between shape pairs, ARAPReg enforces the preservation of the local rigidity in the tangent space specified by the Jacobian. The tangent space is a first-order approximation of the shape space. The smoothness potential ensures that this first-order approximation is accurate, i.e., the rigidity constraint propagates to the shape space's local neighborhood. As we will discuss later, another appealing property of this formulation is that the Jacobian enables us to easily model pose and shape variations (where pose variations are more rigid than shape variations). This goal is hard to achieve using generic pairwise regularizations.

Although the smoothness constraint involves shape pairs, our experiments suggest that there is no need to sample a large number of shape pairs. One interpretation is that deep neural network training has implicit regularizations (c.f. [33]), which promotes smoothness.

4.2. Step II: Jacobian Regularization

We proceed to introduce the local rigidity term r_R that regularizes the Jacobian of the generator. To make the notations uncluttered, we focus on formulating $r_R(g, J)$. Here $g \in \mathcal{R}^{3n}$ denotes a vertex position vector, and $J \in \mathcal{R}^{3n \times k}$ is a Jacobian matrix that specifies infinitesimal perturbations to g .

Our formulation is inspired by the as-rigid-as possible (or ARAP) potential function [43, 49, 55]. This standard model measures the deformation between a pair of shapes. Consider a mesh with vertex position $g \in \mathcal{R}^{3n}$ and the same mesh with perturbed vertex position $g + x \in \mathcal{R}^{3n}$. Denote $O_i \in SO(3)$ as the latent rotation associated with the i -th

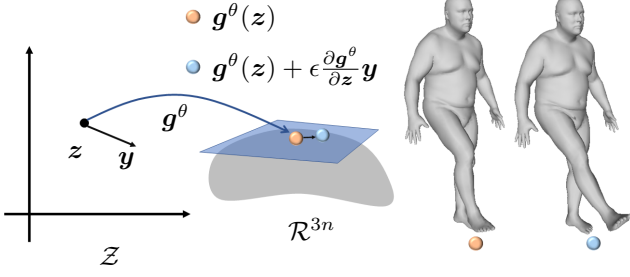


Figure 3: Illustration of configuration for Jacobian regularization. We study infinitesimal deformations incurred by the tangent space at each generated shape $g^\theta(z)$.

vertex. The ARAP deformation between them is

$$f_R(g, x) := \min_{O_i \in SO(3)} \sum_{(i,j) \in \mathcal{E}} \|r_{ij}(O_i, g, x)\|^2 \quad (3)$$

$$r_{ij}(O_i, g, x) := (O_i - I_3)(g_i - g_j) - (x_i - x_j)$$

where $\mathcal{E}$ denotes the edge set of the mesh generator g^θ . Note that we assume $(i, j) \in \mathcal{E}$ if and only if $(j, i) \in \mathcal{E}$.

To introduce a formulation that only depends on the Jacobian of the generator, we consider the Taylor expansion of the as-rigid-as possible potential energy.

Proposition 1 ([20]) *The zero and first-order derivatives of f_R satisfy*

$$f_R(g, 0) = 0, \quad \frac{\partial f_R}{\partial x}(g, 0) = 0.$$

Moreover, the Hessian matrix is given by

$$\frac{\partial^2 f_R}{\partial^2 x}(g, 0) = H_R(g),$$

$$H_R(g) = L \otimes I_3 - A(g)^T D(g)^{-1} A(g), \quad (4)$$

where $L \in \mathcal{R}^{n \times n}$ is the graph Laplacian associated to $\mathcal{E}$; $A(g)$ is a sparse $n \times n$ block matrix; $D(g)$ is a diagonal block matrix. The blocks of $A(g)$ and $D(g)$ are given by

$$A_{ij}(g) = \begin{cases} \sum_{k \in \mathcal{N}(i)} (v_{ik} \times) & i = j \\ -v_{ij} \times & (i, j) \in \mathcal{E} \\ 0 & \text{otherwise} \end{cases}$$

$$D_{ij}(g) = \begin{cases} \sum_{k \in \mathcal{N}(i)} (\|v_{ik}\|^2 I_3 - v_{ik} v_{ik}^T) & i = j \\ 0 & \text{otherwise} \end{cases}$$

where $v_{ij} = g_i - g_j$, and $\mathcal{N}(i)$ collects indices of adjacent vertices of i . Note that H_R is a highly sparse matrix.

Proposition 1 indicates that for each vector $y \in \mathcal{R}^d$ in the parameter space, the ARAP potential between the mesh defined by g and its infinitesimal displacement encoded by $\epsilon J y$ (for a small ϵ) can be approximated as

$$f_R(g, \epsilon J y) \approx \frac{1}{2} \epsilon^2 y^T \bar{H}_R(g, J) y, \quad (5)$$

where $\bar{H}_R(g, J) := J^T H_R(g) J$.

(5) provides the rigidity potential along a direction y in the latent space. Our formulation of r_R seeks to integrate (5) over all possible directions y . To motivate the final formulation of ARAPReg, let us first define an initial potential energy by integrating $y^T \bar{H}_R(g, J) y$ over the unit-sphere S^k in $\mathcal{R}^k$ that specifies all possible y :

$$r_R^{L^2}(g, J) := \frac{k}{\text{Vol}(S^k)} \int_{y \in S^k} y^T \bar{H}_R(g, J) y dy. \quad (6)$$

Proposition 2

$$r_R^{L^2}(g, J) = \text{Tr}(\bar{H}_R(g, J)) = \sum_{i=1}^k \lambda_i(\bar{H}_R(g, J)) \quad (7)$$

where $\lambda_i(\bar{H}_R(g, J))$ is the i -th eigenvalue of $\bar{H}_R(g, J)$.

4.3. Step III: Pose and Shape Variation Modeling

We present a simple formulation that decouples enforcing pose and shape variations. Specifically, the eigenvalues $\lambda_i(\bar{H}_R(g, J)) = u_i^T \bar{H}_R(g, J) u_i$, where u_i is the corresponding eigenvector of $\lambda_i(\bar{H}_R(g, J))$, reveal the deformations in different directions of the tangent space. From the definition of the as-rigid-as possible deformation energy, each vertex's one-ring neighborhood is mostly rigid under pose variations. In contrast, the one-ring neighborhoods may change drastically under shape variations. This means eigenvectors with small eigenvalues correspond to pose variations, while eigenvectors with large eigenvalues correspond to shape variations (See Figure 4).

The limitation of the L2 formulation described in (2) is that all directions are penalized equally. ARAPReg employs a robust norm to model the local rigidity loss to address this issue

$$r_R(g, J) = \sum_{i=1}^k \lambda_i^\alpha(\bar{H}_R(g, J)), \quad (8)$$

where we set $\alpha = \frac{1}{2}$ in this paper. Similar to the effects of using robust norms for outlier removal, (8) imposes small weights on the subspace spanned by eigenvectors of large eigenvalues, which correspond to shape variations. In other words, minimizing (8) minimizes the small eigenvalues of $\bar{H}_R(g, J)$ automatically, which correspond to pose variations. Note that several prior works [63, 11, 2] aimed to decouple pose and shape in the latent space. In contrast, our goal is to model the regularization term by taking pose and shape variations into account.

4.4. Step IV: Final Loss Term

Substituting (6) into (1), we have

$$\mathcal{L}_{reg}(\theta) := \frac{E}{z \sim \mathcal{N}_k} \left(\frac{E}{\delta z \sim s \mathcal{N}_k} \|g^\theta(z + \delta z) - 2g^\theta(z) + g^\theta(z - \delta z)\|^2 + \lambda_R \sum_{i=1}^k \lambda_i^\alpha(\bar{H}_R(g^\theta(z), \frac{\partial g^\theta}{\partial z}(z))) \right) \quad (9)$$



Figure 4: This figure illustrates the local shape space spanned by the eigenvectors of the Hessian of $\bar{H}_R(\mathbf{g}, J)$. The red shape in the center is the reference shape. When moving the latent parameter along the first eigenvector, the shape deformation is locally rigid, exhibiting pose variations (see grey shapes). When moving along the largest eigenvector, the shape deformation possesses local stretching, corresponding to shape variations (see yellow shapes). Finally, when moving along a linear combination of both eigenvectors, the shape exhibits both pose and shape variations (see blue shapes).

In this paper, we set $s = 0.05$ and $\lambda_R = 1$ for all of our experiments.

The major challenge of using (9) for training is to compute the gradient of the Jacobian regularization term. Similar to the formulation of generator smoothness, we introduce a gradient computation approach that only requires computing the derivatives of $\mathbf{g}^\theta$. Please refer to the supp. material for details.

5. Application in Learning Mesh Generators

This section introduces the applications of the unsupervised loss for learning mesh generators. We first introduce the network architecture used in this paper for experimental evaluation. We then introduce how to insert the unsupervised loss $\mathcal{L}_{reg}$ described above into two formulations of training mesh generators, i.e., variational auto-encoders [28, 39, 7] and auto-decoders [59, 61].

5.1. Network Architecture

We focus on describing the decoder network $\mathbf{g}^\theta$. When training variational auto-encoders, we utilize another encoder network $\mathbf{h}^\phi : \mathcal{R}^{3n} \rightarrow \mathcal{Z}$, the mirror of $\mathbf{g}^\theta$ but has different network weights. In other words, $\mathbf{h}^\phi$ has the identical network layers as $\mathbf{g}^\theta$, but the connections are reversed.

In this paper, we model $\mathbf{g}^\theta$ using six layers. The second to the sixth layers are the same as the network architecture of [28]. Motivated from [64], we let the first layer concatenate the latent features associated with each vertex of the coarse mesh as input. Between the input and the first activation is a fully connected layer. Please refer to the supp. material for details.

5.2. Variational Auto-Encoder

Given a collection of training meshes $\mathcal{T} = \{\mathbf{g}_i | 1 \leq i \leq N\}$, we solve the following optimization problem to train the auto-encoder that combines $\mathbf{g}^\theta$ and $\mathbf{h}^\phi$:

$$\min_{\theta, \phi} \frac{1}{N} \sum_{i=1}^N \|\mathbf{g}^\theta(\mathbf{h}^\phi(\mathbf{g}_i)) - \mathbf{g}_i\| + \lambda_{KL} KL(\{\mathbf{h}^\phi(\mathbf{g}_i)\} | \mathcal{N}_k) + \lambda_{reg} \mathcal{L}_{reg}(\theta) \quad (10)$$

where the first two terms of (10) form the standard VAE loss. In this paper, we set $\lambda_{KL} = 1$ and $\lambda_{reg} = 10$. For network training, we employ ADAM [24].

5.3. Auto-Decoder

The auto-decoder formulation [59, 61] replaces the encoder with latent variables $\mathbf{z}_i$ associated with the training meshes:

$$\min_{\theta, \{\mathbf{z}_i\}} \frac{1}{N} \sum_{i=1}^N \|\mathbf{g}^\theta(\mathbf{z}_i) - \mathbf{g}_i\| + \lambda_{KL} KL(\{\mathbf{z}_i\} | \mathcal{N}_k) + \lambda_{reg} \mathcal{L}_{reg}(\theta) \quad (11)$$

where we use the same hyper-parameters as Section 5.2.

We apply alternating minimization to solve (11). The latent parameters $\mathbf{z}_i$ are initialized as an empirical distribution of $\mathcal{N}_k$. When $\mathbf{z}_i$ are fixed, (11) reduces to

$$\min_{\theta} \frac{1}{N} \sum_{i=1}^N \|\mathbf{g}^\theta(\mathbf{z}_i) - \mathbf{g}_i\| + \lambda_{reg} \mathcal{L}_{reg}(\theta) \quad (12)$$

We again employ ADAM [24] to solve (12). Our implementation applies one epoch of optimizing θ for each alternating optimization iteration.

When the network parameters θ are fixed, (11) reduces to

$$\min_{\{z_i\}} \frac{1}{N} \sum_{i=1}^N \|g^\theta(z_i) - g_i\| + \lambda_{KL} KL(\{z_i\} | \mathcal{N}_k) \quad (13)$$

We again employ ADAM [24] to optimize z_i . Similarly, our implementation applies one epoch of optimizing z_i for each alternating iteration. The total number of alternating iterations is set as 30 in this paper.

6. Experimental Evaluation

This section presents an experimental evaluation of ARAPReg. In Section 6.1, we present the experimental setup. We then analyze the experimental results in Section 6.2. Finally, Section 6.3 and Section 6.4 describe an ablation study of the ASARReg loss and an evaluation of the shape interpolation application. Due to space issues, we defer more results and comparisons to the supp. material.

6.1. Experimental Setup

Datasets. The experimental evaluation considers three datasets: **DFAUST** [5], **SMAL** [66], and **Bone**. The DFAUST dataset consists of 37,197 human shapes for training and 4,264 shapes for testing. All the shapes are generated using the SMPL model [29]. For the SMAL dataset, we randomly generate 400 shapes for following the shape sampling method in [15], where latent vectors are sampled from a normal distribution with zero mean and 0.2 standard deviations. We split them into 300 training shapes and 100 testing shapes. The Bone dataset consists of four categories of real bones: Femur, Tibia, Pelvis, and Scapula, where each category has 40 training and 10 testing shapes. The consistent correspondences are obtained from interpolating landmark correspondences marked by experts.

Baseline approaches. We evaluate on four baselines: **SP-Disentangle** [63], **CoMA** [39], **3DMM** [7], and **MeshConv** [64]. They together represent the state-of-the-art results on learning mesh generators from a collection of meshes with dense correspondences. We evaluate the effectiveness of ARAPReg on these baselines and the absolute performance of our approach against these baselines.

Evaluation metrics. Besides qualitative evaluations, we employ the reconstruction error metric (c.f. [39, 7, 64]) for quantitative evaluations. Specifically, we compute the average per-vertex Euclidean distance for the input and reconstructed meshes. For VAE, the latent variable is given by the encoder. For AD, we optimize the latent variable to find the best reconstruction (c.f. [59, 61]). The output shape is obtained by feeding the latent variable to the decoder.

6.2. Analysis of Results

Table 1 compares our approach and baseline approaches in terms of the reconstruction error. Under the AD frame-

	DFAUST	SMAL	Bone
SP-Disentangle. [63]	10.02	21.32	5.34
COMA[39]	8.80	14.52	4.14
3DMM[7]	7.39	17.78	4.03
MeshConv[64]	5.43	8.01	4.47
Ours-VAE (L1)	5.45	9.11	4.09
Ours-VAE (L1 + ARAP)	4.87	7.82	3.85
Ours-AD (L1)	5.17	8.74	3.91
Ours-AD (L1 + ARAP)	4.52	6.68	3.76

Table 1: Correspondence-based MSE reconstruction error (mm) on test sets of Dfaust, SMAL and Bone.

	3DMM [7]	COMA [39]	MeshConv [64]
No-Reg	7.39	8.80	5.43
ARAPReg	6.72	4.87	5.02

Table 2: The effects of ARAPReg on different baselines on DFAUST dataset. The first row shows reported MSE reconstruction errors in their papers, and the second row shows results with ARAPReg on the same architecture with VAE training. ARAPReg achieves improvements on various baselines with different architectures.

work, our approach reduces the reconstruction error of baseline approaches by 16.8%, 16.6%, and 6.7% on DFAUST, SMAL, and Bone, respectively. As the optimal latent-variable is optimized, AD framework achieves better quality than VAE framework.

Figure 5 illustrates the reconstruction errors visually. Our approach improves from baseline approaches considerably. In particular, it improves from the top-performing approach MeshConv [64] at locations with large deformations (e.g., arms of humans) and non-rigid deformations (e.g., arms and torsos of humans). These improvements come from modeling the preservation of the local rigidity among neighboring shapes in the underlying shape space. Please refer to the supp. material for more results.

Figure 8 and the supp. material shows randomly generated shapes under our trained full VAE and AD models (i.e., with ARAPReg). We can see that the generated shapes nicely preserve important shape features such as fingers and faces of human shapes and tails of animal shapes. Moreover, the generated shapes are different from the closest training shape, indicating that the learned mesh generator has a strong generalization ability.

6.3. Ablation Study

Table 1 shows our quantitative reconstruction results with and without ARAPReg under the VAE and AD settings. The effects of ARAPReg are salient. Under the AD setting, ARAPReg leads to 12.6%, 23.5%, and 4.5% reductions of the reconstruction error on DFAUST, SMAL, Bone, respectively. Table 2 further shows the effect of ARAPReg on various baselines under the VAE reconstruction pipeline. ARAPReg reduces the reconstruction error by building a better shape space that preserves the local rigidity constraint.

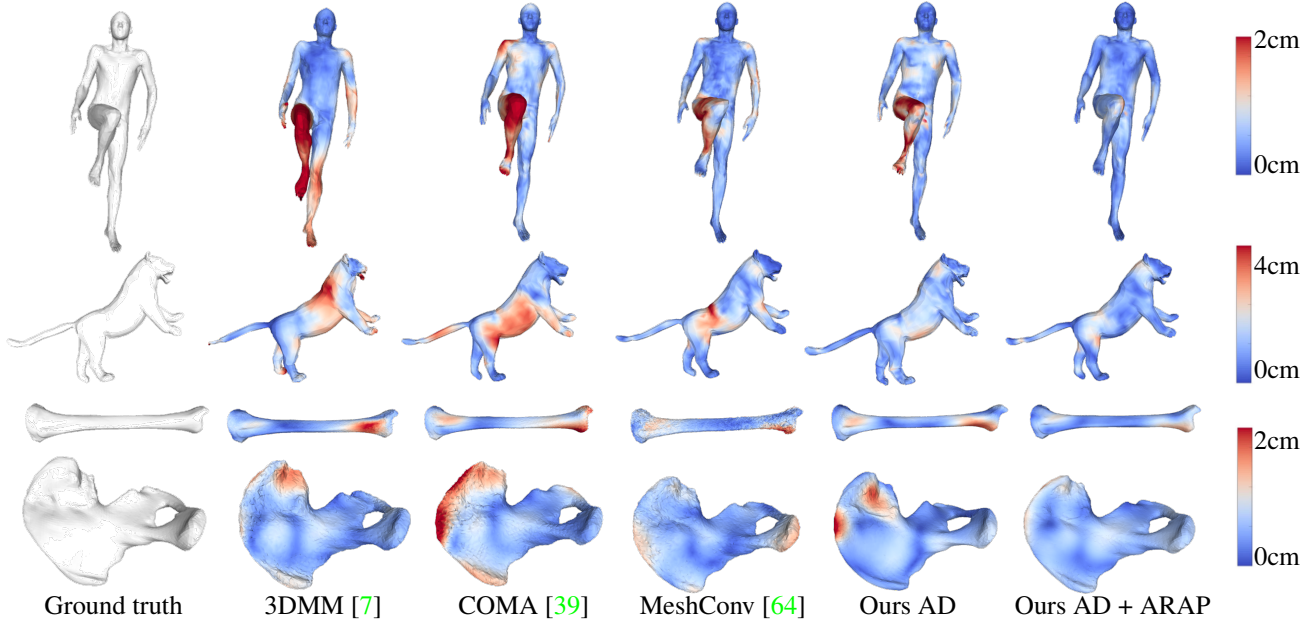


Figure 5: Qualitative comparison of reconstruction results. We show results using the AD generator w/w.o ARAPReg. Compared with baseline approaches, ours results with ARAPReg present less distortions and are locally smoother.

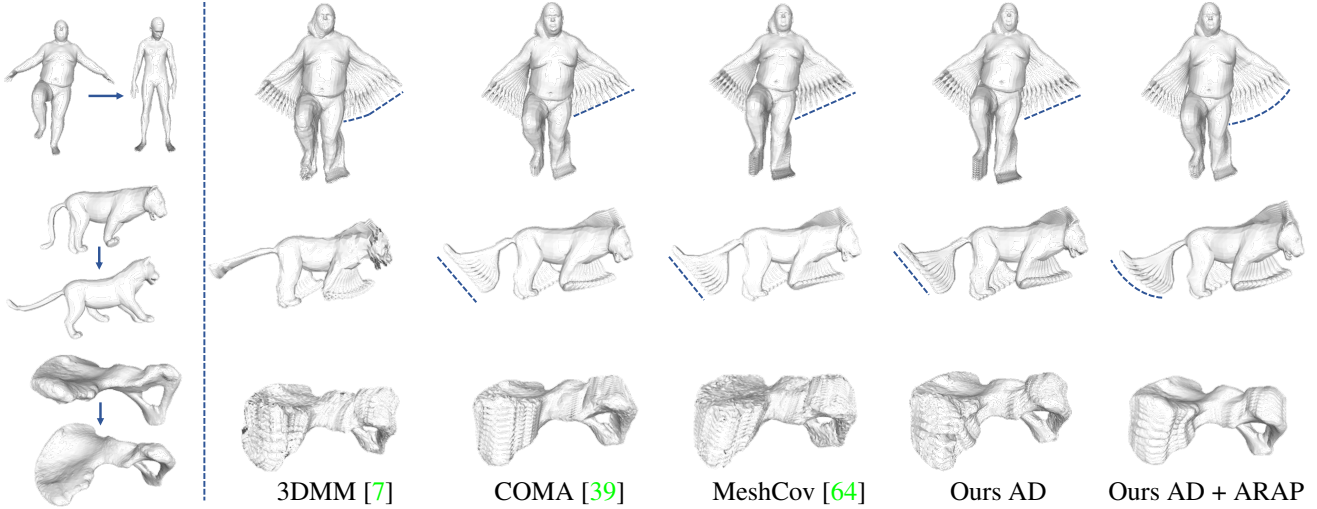


Figure 6: Interpolation results. The left column shows three groups of source and target shapes (connected by blue arrows). The remaining columns show ten intermediate shapes by linearly interpolating the latent codes of the source and target shapes. We show results using the AD generator w/w.o ARAPReg. Compared with baseline approaches, our results with ARAPReg show much smoother and more shape-preserving deformations.

6.4. Shape Interpolation

We proceed to evaluate the effects of ARAPReg for the application of shape interpolation. Given two shapes g_1 and g_2 , we first obtain their corresponding latent parameters z_1 and z_2 . For the VAE model, z_i comes from the encoder. For the AD model, z_i comes from optimizing the reconstruction error. The interpolation is then done by linearly interpolating z_1 and z_2 .

Figure 6 compares our approach and baseline approaches on shape interpolation. We can see that our approach’s in-

terpolated shapes are smoother and more shape-preserving than those of the baseline approaches. Specifically, prominent shape features such as fingers are better preserved in our approach. Moreover, our approach introduces less distortion among joint regions.

6.5. Shape Extrapolation

We also evaluate the effects of ARAPReg for the application of shape extrapolation. Given a center shape g , we first obtain its corresponding latent parameters z . For the VAE model, z comes from the encoder. For the AD

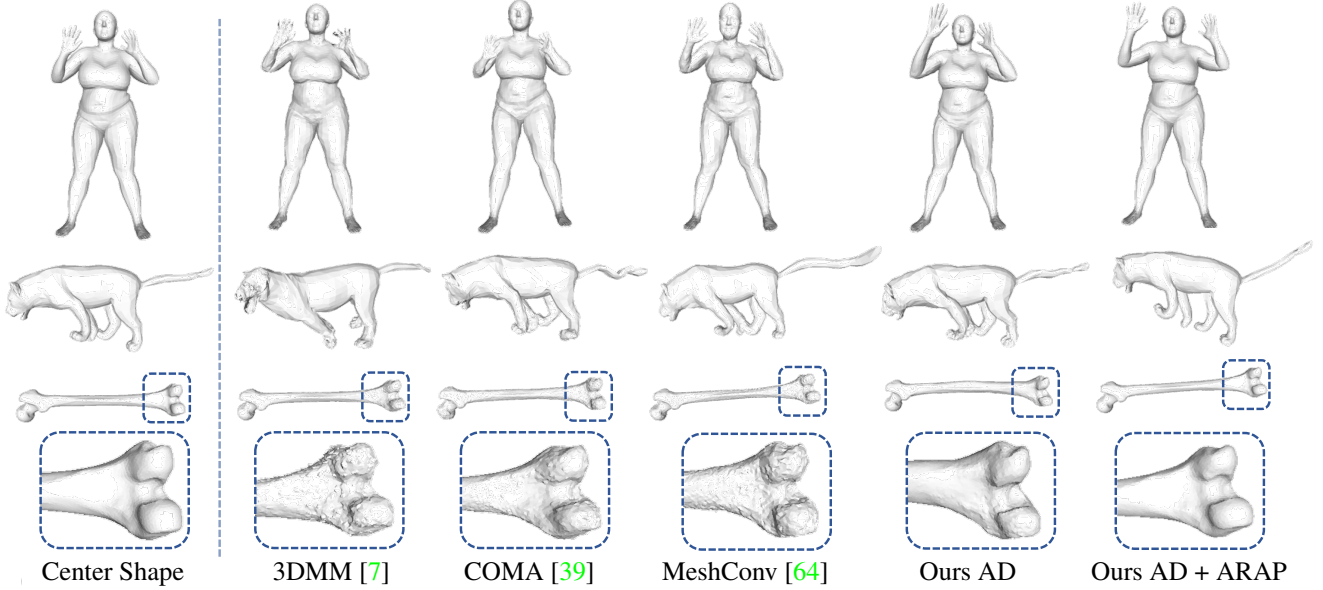
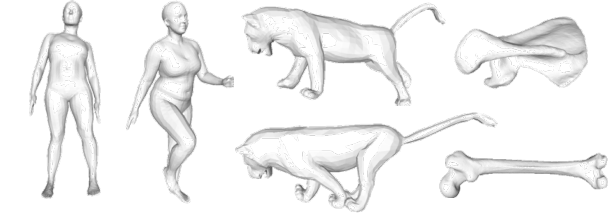


Figure 7: Extrapolation results. Around one test center shape (left column), we randomly perturb its latent code z within an Euclidean ball to generate perturbed shapes. We show results using the AD generator w/o ARAPReg. Our results with ARAPReg exhibit smooth and feature-preserving deformations.

Generated Shapes



Closest Shapes

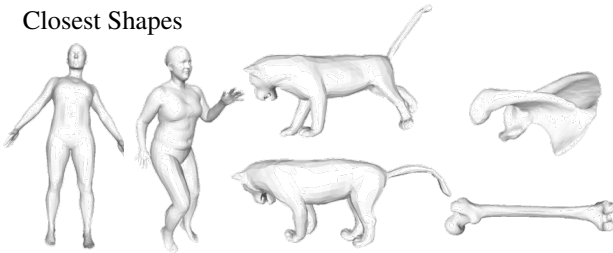


Figure 8: Randomly generated shapes from our AD + ARAPReg framework and their closest shapes in the training set. Our network is able to generate reasonable shapes that are not in the training shape collection. Due to space constraints, results of our VAE framework are in the supp. material.

model, z comes from optimizing the reconstruction error. The extrapolation is then done by randomly sampling $\tilde{z} \sim z + \mathcal{N}(0, \sigma^2 S)$, where S denotes scale for each latent dimension. We choose $\sigma = 0.2$ for all datasets.

Figure 7 compares our approach and baseline approaches on shape extrapolation. We can see that our approach’s generated shapes are smoother and more reasonable than base-

line approaches in areas such as tails of animals, hands and arms of human.

7. Conclusions and Limitations

This paper introduces ARAPReg, an unsupervised loss functional for training shape generators. Experimental results show that enforcing this loss on meshed shape generators improves their performance. The resulting mesh generators produce novel generated shapes that are shape-preserving at multiple scales.

ARAPReg has several limitations which can inspire future work. First, so far, ARAPReg only applies to training datasets with given correspondences. An interesting problem is to address unorganized shape collections that do not possess dense correspondences. Besides pre-computing correspondences, a promising direction is to explore the simultaneous learning of the shape correspondences and the shape generator. Another limitation of ARAPReg is that it targets realistically deformable shapes. Future directions are to study how to extend the formulation to handle synthetically generated shapes of any form and function.

Acknowledgement. Chandrajit Bajaj would like to acknowledge the support from NIH-R01GM117594, by the Peter O’Donnell Foundation, and in part from a grant from the Army Research Office accomplished under Cooperative Agreement Number W911NF-19-2-0333. Junfeng Jiang would like to acknowledge support from Jiangsu NSF under Grant BK20181158 and NSF of China (NSFC) under Grant 61772172. Qixing Huang is supported by NSF Career IIS-2047677, NSF HDR TRIPODS-1934932, and Gifts from Wormpex AI Research and Snap Inc.

References

- [1] Panos Achlioptas, Olga Diamanti, Ioannis Mitliagkas, and Leonidas J. Guibas. Learning representations and generative models for 3d point clouds. In Jennifer G. Dy and Andreas Krause, editors, *Proceedings of the 35th International Conference on Machine Learning, ICML 2018, Stockholm, Sweden, July 10-15, 2018*, volume 80 of *Proceedings of Machine Learning Research*, pages 40–49, Stockholm, Sweden, 2018. PMLR.
- [2] Tristan Aumentado-Armstrong, Stavros Tsogkas, Allan D. Jepson, and Sven J. Dickinson. Geometric disentanglement for generative latent shape models. In *2019 IEEE/CVF International Conference on Computer Vision, ICCV 2019, Seoul, Korea (South), October 27 - November 2, 2019*, pages 8180–8189. IEEE, 2019.
- [3] Elena Balashova, Vivek Singh, Jiangping Wang, Brian Teixeira, Terrence Chen, and Thomas A. Funkhouser. Structure-aware shape synthesis. In *2018 International Conference on 3D Vision, 3DV 2018, Verona, Italy, September 5-8, 2018*, pages 140–149. IEEE Computer Society, 2018.
- [4] Heli Ben-Hamu, Haggai Maron, Itay Kezurer, Gal Avineri, and Yaron Lipman. Multi-chart generative surface modeling. *ACM Trans. Graph.*, 37(6):215:1–215:15, 2018.
- [5] Federica Bogo, Javier Romero, Gerard Pons-Moll, and Michael J. Black. Dynamic FAUST: Registering human bodies in motion. In *IEEE Conf. on Computer Vision and Pattern Recognition (CVPR)*, July 2017.
- [6] Mario Botsch and Olga Sorkine. On linear variational surface deformation methods. *IEEE Transactions on Visualization and Computer Graphics*, 14(1):213–230, 2008.
- [7] Giorgos Bouritsas, Sergiy Bokhnyak, Stylianos Ploumpis, Michael Bronstein, and Stefanos Zafeiriou. Neural 3d morphable models: Spiral convolutional networks for 3d shape representation learning and generation. In *The IEEE International Conference on Computer Vision (ICCV)*, 2019.
- [8] Zhangjie Cao, Qixing Huang, and Karthik Ramani. 3d object classification via spherical projections. In *2017 International Conference on 3D Vision, 3DV 2017, Qingdao, China, October 10-12, 2017*, pages 566–574, Qingdao, China, 2017. IEEE Computer Society.
- [9] Zhiqin Chen and Hao Zhang. Learning implicit fields for generative shape modeling. In *IEEE Conference on Computer Vision and Pattern Recognition, CVPR 2019, Long Beach, CA, USA, June 16-20, 2019*, pages 5939–5948, Long Beach, CA, USA, 2019. Computer Vision Foundation / IEEE.
- [10] Taco S. Cohen, Mario Geiger, Jonas Köhler, and Max Welling. Spherical cnns. In *6th International Conference on Learning Representations, ICLR 2018, Vancouver, BC, Canada, April 30 - May 3, 2018, Conference Track Proceedings*. OpenReview.net, 2018.
- [11] Luca Cosmo, Antonio Norelli, Oshri Halimi, Ron Kimmel, and Emanuele Rodolà. LIMP: learning latent shape representations with metric preservation priors. In Andrea Vedaldi, Horst Bischof, Thomas Brox, and Jan-Michael Frahm, editors, *Computer Vision - ECCV 2020 - 16th European Conference, Glasgow, UK, August 23-28, 2020, Proceedings, Part III*, volume 12348 of *Lecture Notes in Computer Science*, pages 19–35. Springer, 2020.
- [12] Angela Dai, Charles Ruizhongtai Qi, and Matthias Nießner. Shape completion using 3d-encoder-predictor cnns and shape synthesis. In *2017 IEEE Conference on Computer Vision and Pattern Recognition, CVPR 2017, Honolulu, HI, USA, July 21-26, 2017*, pages 6545–6554, Honolulu, HI, USA, 2017. IEEE Computer Society.
- [13] Carlos Esteves, Christine Allen-Blanchette, Ameesh Makadia, and Kostas Daniilidis. 3d object classification and retrieval with spherical cnns. *CoRR*, abs/1711.06721:1–9, 2017.
- [14] Lin Gao, Jie Yang, Tong Wu, Yu-Jie Yuan, Hongbo Fu, Yu-Kun Lai, and Hao Zhang. Sdm-net: Deep generative network for structured deformable mesh. *ACM Trans. Graph.*, 38(6), Nov. 2019.
- [15] Thibault Groueix, Matthew Fisher, Vladimir G. Kim, Bryan C. Russell, and Mathieu Aubry. 3d-coded : 3d correspondences by deep deformation. In *In ECCV 2018*, 2018.
- [16] Thibault Groueix, Matthew Fisher, Vladimir G. Kim, Bryan C. Russell, and Mathieu Aubry. A papier-mâché approach to learning 3d surface generation. In *2018 IEEE Conference on Computer Vision and Pattern Recognition, CVPR 2018, Salt Lake City, UT, USA, June 18-22, 2018*, pages 216–224, 2018.
- [17] Marc Habermann, Weipeng Xu, Michael Zollhöfer, Gerard Pons-Moll, and Christian Theobalt. Deepcap: Monocular human performance capture using weak supervision. In *2020 IEEE/CVF Conference on Computer Vision and Pattern Recognition, CVPR 2020, Seattle, WA, USA, June 13-19, 2020*, pages 5051–5062. IEEE, 2020.
- [18] Rana Hanocka, Amir Hertz, Noa Fish, Raja Giryes, Shachar Fleishman, and Daniel Cohen-Or. Meshcnn: A network with an edge. *ACM Trans. Graph.*, 38(4), July 2019.
- [19] Kaiming He, Georgia Gkioxari, Piotr Dollár, and Ross B. Girshick. Mask R-CNN. *IEEE Trans. Pattern Anal. Mach. Intell.*, 42(2):386–397, 2020.
- [20] Qi-Xing Huang, Martin Wicke, Bart Adams, and Leonidas J. Guibas. Shape decomposition using modal analysis. *Comput. Graph. Forum*, 28(2):407–416, 2009.
- [21] Angjoo Kanazawa, Michael J. Black, David W. Jacobs, and Jitendra Malik. End-to-end recovery of human shape and pose. In *2018 IEEE Conference on Computer Vision and Pattern Recognition, CVPR 2018, Salt Lake City, UT, USA, June 18-22, 2018*, pages 7122–7131. IEEE Computer Society, 2018.
- [22] Hiroharu Kato, Yoshitaka Ushiku, and Tatsuya Harada. Neural 3d mesh renderer. In *2018 IEEE Conference on Computer Vision and Pattern Recognition, CVPR 2018, Salt Lake City, UT, USA, June 18-22, 2018*, pages 3907–3916, Salt Lake City, UT, USA, 2018. IEEE Computer Society.
- [23] Martin Kilian, Niloy J. Mitra, and Helmut Pottmann. Geometric modeling in shape space. *ACM Trans. Graph.*, 26(3):64–es, July 2007.
- [24] Diederik P. Kingma and Jimmy Ba. Adam: A method for stochastic optimization. In Yoshua Bengio and Yann LeCun, editors, *3rd International Conference on Learning Representations, ICLR 2015, San Diego, CA, USA, May 7-9, 2015, Conference Track Proceedings*, 2015.
- [25] Chun-Liang Li, Manzil Zaheer, Yang Zhang, Barnabás Póczos, and Ruslan Salakhutdinov. Point cloud GAN. In *Deep Generative Models for Highly Structured Data, ICLR 2019 Workshop, New Orleans, Louisiana, United States, May 6, 2019*, pages 1–19, New Orleans, Louisiana, United States, 2019. OpenReview.net.

- [26] Jun Li, Kai Xu, Siddhartha Chaudhuri, Ersin Yumer, Hao Zhang, and Leonidas Guibas. Grass: Generative recursive autoencoders for shape structures. *ACM Trans. Graph.*, 36(4):52:1–52:14, July 2017.
- [27] Xueting Li, Sifei Liu, Shalini De Mello, Kihwan Kim, Xiaolong Wang, Ming-Hsuan Yang, and Jan Kautz. Online adaptation for consistent mesh reconstruction in the wild. In Hugo Larochelle, Marc’Aurelio Ranzato, Raia Hadsell, Maria-Florina Balcan, and Hsuan-Tien Lin, editors, *Advances in Neural Information Processing Systems 33: Annual Conference on Neural Information Processing Systems 2020, NeurIPS 2020, December 6-12, 2020, virtual*, 2020.
- [28] Or Litany, Alexander M. Bronstein, Michael M. Bronstein, and Ameesh Makadia. Deformable shape completion with graph convolutional autoencoders. In *2018 IEEE Conference on Computer Vision and Pattern Recognition, CVPR 2018, Salt Lake City, UT, USA, June 18-22, 2018*, pages 1886–1895, Salt Lake City, UT, USA, 2018. IEEE Computer Society.
- [29] Matthew Loper, Naureen Mahmood, Javier Romero, Gerard Pons-Moll, and Michael J. Black. Smpl: A skinned multi-person linear model. *ACM Trans. Graph.*, 34(6), Oct. 2015.
- [30] Zhaoliang Lun, Matheus Gadelha, Evangelos Kalogerakis, Subhransu Maji, and Rui Wang. 3d shape reconstruction from sketches via multi-view convolutional networks. In *2017 International Conference on 3D Vision, 3DV 2017, Qingdao, China, October 10-12, 2017*, pages 67–77, Qingdao, China, 2017. IEEE Computer Society.
- [31] Haggai Maron, Meirav Galun, Noam Aigerman, Miri Trope, Nadav Dym, Ersin Yumer, Vladimir G. Kim, and Yaron Lipman. Convolutional neural networks on surfaces via seamless toric covers. *ACM Trans. Graph.*, 36(4), July 2017.
- [32] Lars M. Mescheder, Michael Oechsle, Michael Niemeyer, Sebastian Nowozin, and Andreas Geiger. Occupancy networks: Learning 3d reconstruction in function space. In *IEEE Conference on Computer Vision and Pattern Recognition, CVPR 2019, Long Beach, CA, USA, June 16-20, 2019*, pages 4460–4470, Long Beach, CA, USA, 2019. Computer Vision Foundation / IEEE.
- [33] Behnam Neyshabur, Ryota Tomioka, Ruslan Salakhutdinov, and Nathan Srebro. Geometry of optimization and implicit regularization in deep learning. *CoRR*, abs/1705.03071, 2017.
- [34] Eunbyung Park, Jimei Yang, Ersin Yumer, Duygu Ceylan, and Alexander C. Berg. Transformation-grounded image generation network for novel 3d view synthesis. In *2017 IEEE Conference on Computer Vision and Pattern Recognition, CVPR 2017, Honolulu, HI, USA, July 21-26, 2017*, pages 702–711, Honolulu, HI, USA, 2017. IEEE Computer Society.
- [35] Jeong Joon Park, Peter Florence, Julian Straub, Richard A. Newcombe, and Steven Lovegrove. DeepSDF: Learning continuous signed distance functions for shape representation. In *IEEE Conference on Computer Vision and Pattern Recognition, CVPR 2019, Long Beach, CA, USA, June 16-20, 2019*, pages 165–174. Computer Vision Foundation / IEEE, 2019.
- [36] William S. Peebles, John Peebles, Jun-Yan Zhu, Alexei A. Efros, and Antonio Torralba. The hessian penalty: A weak prior for unsupervised disentanglement. In Andrea Vedaldi, Horst Bischof, Thomas Brox, and Jan-Michael Frahm, editors, *Computer Vision - ECCV 2020 - 16th European Conference, Glasgow, UK, August 23-28, 2020, Proceedings, Part VI*, volume 12351 of *Lecture Notes in Computer Science*, pages 581–597. Springer, 2020.
- [37] Helmut Pottmann, Qixing Huang, Bailin Deng, Alexander Schiftner, Martin Kilian, Leonidas Guibas, and Johannes Wallner. Geodesic patterns. *ACM Trans. Graph.*, 29(4), July 2010.
- [38] Marie-Julie Rakotosaona and Maks Ovsjanikov. Intrinsic point cloud interpolation via dual latent space navigation. In Andrea Vedaldi, Horst Bischof, Thomas Brox, and Jan-Michael Frahm, editors, *Computer Vision - ECCV 2020 - 16th European Conference, Glasgow, UK, August 23-28, 2020, Proceedings, Part II*, volume 12347 of *Lecture Notes in Computer Science*, pages 655–672. Springer, 2020.
- [39] Anurag Ranjan, Timo Bolkart, Soubhik Sanyal, and Michael J Black. Generating 3d faces using convolutional mesh autoencoders. In *Proceedings of the European Conference on Computer Vision (ECCV)*, pages 704–720, 2018.
- [40] Daniel Ritchie, Anna Thomas, Pat Hanrahan, and Noah D. Goodman. Neurally-guided procedural models: Amortized inference for procedural graphics programs using neural networks. In *Proceedings of the 30th International Conference on Neural Information Processing Systems, NIPS’16*, pages 622–630, USA, 2016. Curran Associates Inc.
- [41] Gopal Sharma, Rishabh Goyal, Difan Liu, Evangelos Kalogerakis, and Subhransu Maji. Csgnet: Neural shape parser for constructive solid geometry. In *2018 IEEE Conference on Computer Vision and Pattern Recognition, CVPR 2018, Salt Lake City, UT, USA, June 18-22, 2018*, pages 5515–5523, Los Alamitos, CA, 2018. IEEE Computer Society.
- [42] Ayan Sinha, Asim Unmesh, Qixing Huang, and Karthik Ramani. Surfnet: Generating 3d shape surfaces using deep residual networks. In *2017 IEEE Conference on Computer Vision and Pattern Recognition, CVPR 2017, Honolulu, HI, USA, July 21-26, 2017*, pages 791–800, 2017.
- [43] Olga Sorkine and Marc Alexa. As-rigid-as-possible surface modeling. In *Proceedings of the Fifth Eurographics Symposium on Geometry Processing, SGP ’07*, page 109–116, Goslar, DEU, 2007. Eurographics Association.
- [44] Qingyang Tan, Lin Gao, Yu-Kun Lai, and Shihong Xia. Variational autoencoders for deforming 3d mesh models. In *2018 IEEE Conference on Computer Vision and Pattern Recognition, CVPR 2018, Salt Lake City, UT, USA, June 18-22, 2018*, pages 5841–5850, Salt Lake City, UT, USA, 2018. IEEE Computer Society.
- [45] Maxim Tatarchenko, Alexey Dosovitskiy, and Thomas Brox. Octree generating networks: Efficient convolutional architectures for high-resolution 3d outputs. In *IEEE International Conference on Computer Vision, ICCV 2017, Venice, Italy, October 22-29, 2017*, pages 2107–2115, Venice, Italy, 2017. IEEE Computer Society.
- [46] Edgar Tretschk, Ayush Tewari, Michael Zollhöfer, Vladislav Golyanik, and Christian Theobalt. DEMEA: deep mesh autoencoders for non-rigidly deforming objects. In Andrea Vedaldi, Horst Bischof, Thomas Brox, and Jan-Michael Frahm, editors, *Computer Vision - ECCV 2020 - 16th European Conference, Glasgow, UK, August 23-28, 2020, Proceedings, Part IV*, volume 12349 of *Lecture Notes in Computer Science*, pages 601–617, Glasgow, UK, 2020. Springer.
- [47] Shubham Tulsiani, Hao Su, Leonidas J. Guibas, Alexei A. Efros, and Jitendra Malik. Learning shape abstractions by assembling volumetric primitives. In *The IEEE Conference*

on *Computer Vision and Pattern Recognition (CVPR)*, July 2017.

- [48] Nitika Verma, Edmond Boyer, and Jakob Verbeek. Feastnet: Feature-steered graph convolutions for 3d shape analysis. In *2018 IEEE Conference on Computer Vision and Pattern Recognition, CVPR 2018, Salt Lake City, UT, USA, June 18-22, 2018*, pages 2598–2606, Salt Lake City, UT, USA, 2018. IEEE Computer Society.
- [49] Michael Wand, Philipp Jenke, Qi-Xing Huang, Martin Bokeloh, Leonidas J. Guibas, and Andreas Schilling. Reconstruction of deforming geometry from time-varying point clouds. In Alexander G. Belyaev and Michael Garland, editors, *Proceedings of the Fifth Eurographics Symposium on Geometry Processing, Barcelona, Spain, July 4-6, 2007*, volume 257 of *ACM International Conference Proceeding Series*, pages 49–58, Avenue de Frontenex 32, 1207 Geneve, Switzerland, 2007. Eurographics Association.
- [50] Nanyang Wang, Yinda Zhang, Zhuwen Li, Yanwei Fu, Wei Liu, and Yu-Gang Jiang. Pixel2mesh: Generating 3d mesh models from single RGB images. In *Computer Vision - ECCV 2018 - 15th European Conference, Munich, Germany, September 8-14, 2018, Proceedings, Part XI*, pages 55–71, 2018.
- [51] Yifan Wang, Noam Aigerman, Vladimir G. Kim, Siddhartha Chaudhuri, and Olga Sorkine-Hornung. Neural cages for detail-preserving 3d deformations. In *2020 IEEE/CVF Conference on Computer Vision and Pattern Recognition, CVPR 2020, Seattle, WA, USA, June 13-19, 2020*, pages 72–80, Seattle, WA, USA, 2020. IEEE.
- [52] Jiajun Wu, Tianfan Xue, Joseph J. Lim, Yuandong Tian, Joshua B. Tenenbaum, Antonio Torralba, and William T. Freeman. 3d interpreter networks for viewer-centered wireframe modeling. *Int. J. Comput. Vis.*, 126(9):1009–1026, 2018.
- [53] Jiajun Wu, Chengkai Zhang, Tianfan Xue, Bill Freeman, and Josh Tenenbaum. Learning a probabilistic latent space of object shapes via 3d generative-adversarial modeling. In Daniel D. Lee, Masashi Sugiyama, Ulrike von Luxburg, Isabelle Guyon, and Roman Garnett, editors, *Advances in Neural Information Processing Systems 29: Annual Conference on Neural Information Processing Systems 2016, December 5-10, 2016, Barcelona, Spain*, pages 82–90. OpenReview.net, Barcelona, Spain, 2016.
- [54] Kai Xu, Vladimir G. Kim, Qixing Huang, Niloy Mitra, and Evangelos Kalogerakis. Data-driven shape analysis and processing. In *SIGGRAPH ASIA 2016 Courses*, SA '16, New York, NY, USA, 2016. Association for Computing Machinery.
- [55] Weiwei Xu, Kun Zhou, Yizhou Yu, Qifeng Tan, Qunsheng Peng, and Baining Guo. Gradient domain editing of deforming mesh sequences. In *ACM SIGGRAPH 2007 Papers*, SIGGRAPH '07, page 84–es, New York, NY, USA, 2007. Association for Computing Machinery.
- [56] Guandao Yang, Xun Huang, Zekun Hao, Ming-Yu Liu, Serge J. Belongie, and Bharath Hariharan. Pointflow: 3d point cloud generation with continuous normalizing flows. In *2019 IEEE/CVF International Conference on Computer Vision, ICCV 2019, Seoul, Korea (South), October 27 - November 2, 2019*, pages 4540–4549, Seoul, Korea (South), 2019. IEEE.
- [57] Yaoqing Yang, Chen Feng, Yiru Shen, and Dong Tian. Foldingnet: Point cloud auto-encoder via deep grid deformation. In *2018 IEEE Conference on Computer Vision and Pattern Recognition, CVPR 2018, Salt Lake City, UT, USA, June 18-22, 2018*, pages 206–215, Salt Lake City, UT, USA, 2018. IEEE Computer Society.
- [58] Yong-Liang Yang, Yi-Jun Yang, Helmut Pottmann, and Niloy J. Mitra. Shape space exploration of constrained meshes. In *Proceedings of the 2011 SIGGRAPH Asia Conference, SA '11*, New York, NY, USA, 2011. Association for Computing Machinery.
- [59] Zhenpei Yang, Lihang Liu, and Qixing Huang. Learning generative neural networks for 3d colorization. In *Proceedings of the Thirty-Second AAAI Conference on Artificial Intelligence, (AAAI-18), the 30th innovative Applications of Artificial Intelligence (IAAI-18), and the 8th AAAI Symposium on Educational Advances in Artificial Intelligence (EAAI-18), New Orleans, Louisiana, USA, February 2-7, 2018*, pages 2580–2587, New York, USA, 2018. AAAI.
- [60] Mehmet Ersin Yümer and Niloy J. Mitra. Learning semantic deformation flows with 3d convolutional networks. In Bastian Leibe, Jiri Matas, Nicu Sebe, and Max Welling, editors, *Computer Vision - ECCV 2016 - 14th European Conference, Amsterdam, The Netherlands, October 11-14, 2016, Proceedings, Part VI*, volume 9910 of *Lecture Notes in Computer Science*, pages 294–311, Amsterdam, The Netherlands, 2016. Springer.
- [61] Amir Zadeh, Yao-Chong Lim, Paul Pu Liang, and Louis-Philippe Morency. Variational auto-decoder. *arXiv preprint arXiv:1903.00840*, 2019.
- [62] Zaiwei Zhang, Zhenpei Yang, Congyang Ma, Linjie Luo, Alexander Huth, Etienne Vouga, and Qixing Huang. Deep generative modeling for scene synthesis via hybrid representations. *ACM Transaction on Graphics*, 2, 2020.
- [63] Keyang Zhou, Bharat Lal Bhatnagar, and Gerard Pons-Moll. Unsupervised shape and pose disentanglement for 3d meshes. In Andrea Vedaldi, Horst Bischof, Thomas Brox, and Jan-Michael Frahm, editors, *Computer Vision - ECCV 2020 - 16th European Conference, Glasgow, UK, August 23-28, 2020, Proceedings, Part XXII*, volume 12367 of *Lecture Notes in Computer Science*, pages 341–357. Springer, 2020.
- [64] Yi Zhou, Chenglei Wu, Zimo Li, Chen Cao, Yuting Ye, Jason Saragih, Hao Li, and Yaser Sheikh. Fully convolutional mesh autoencoder using efficient spatially varying kernels. *arXiv preprint arXiv:2006.04325*, 2020.
- [65] Chuhang Zou, Ersin Yumer, Jimei Yang, Duygu Ceylan, and Derek Hoiem. 3d-prnn: Generating shape primitives with recurrent neural networks. In *IEEE International Conference on Computer Vision, ICCV 2017, Venice, Italy, October 22-29, 2017*, pages 900–909, Venice, Italy, 2017. IEEE Computer Society.
- [66] Silvia Zuffi, Angjoo Kanazawa, David W Jacobs, and Michael J Black. 3d menagerie: Modeling the 3d shape and pose of animals. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pages 6365–6373, 2017.

	DFAUST	SMAL	Bone
W.o. Decoupling	4.90	7.23	3.82
With Decoupling	4.52	6.68	3.76

Table 3: Ablation study on shape and pose variation. In w.o. decoupling setting, all directions are penalized equally. With decoupling setting is the setting in the main paper, where pose directions are penalized more than shape directions.

	DFAUST	SMAL
No ARAP	5.17	8.74
ARAP Deform.	24.55	7.67
Ours	4.52	6.68

Table 4: Comparison between our method and the traditional ARAP deformation method. We show reconstruction errors of AD model without ARAP, with traditional ARAP and with our method. The traditional method couldn't handle large pose variation and shape distortion.

A. More Quantitative Results

A.1. Ablation Study on Pose and Shape Variation in Section 4.3

In the Section 4.2, we introduced decoupling shape and pose variations to improve ARAPReg. Here we show an ablation study of this decoupling. In Table.3, we show MSE reconstruction error in AD framework w/w.o shape and pose decoupling. Specifically, in the non-decoupling setting, we use the L2 formulation in Proposition 2, where all directions are penalized equally.

A.2. Comparison with ARAP deformation from the base mesh

Here we show the comparison between our method and the traditional ARAP deformation method, where an ARAP deformation is applied between the base mesh and the output mesh for regularization (c.f. [17, 27, 63]). In Table 4, we show results on DFAUST and SMAL datasets. On DFAUST dataset, there are large deformations among the underlying shapes, and the approach of enforcing an ARAP loss to the base shape is significantly worse than without the ARAP loss. In the SMAL dataset, we pick all samples with the same shape but different poses, the ARAP loss to the base shape offers slight performance gains. However, ARAPReg still outperforms this simple baseline considerably.

B. More Implementation Details

B.1. Model Architecture

Our VAE model consists of a shape encoder and a decoder. Our AD model only contains a decoder. Both encoder and decoder are composed of Chebyshev convolutional filters with $K = 6$ Chebyshev polynomials [39]. The VAE model architecture is based on [39]. We sample 4 resolutions of the mesh connections of the template mesh.

The encoder is stacked by 4 blocks of convolution + down-sampling layers. The decoder is stacked by 4 blocks of convolution + up-sampling layers. There's two fully connected layers connecting the encoder, latent variable and the decoder. For the full details, please refer to our Github repository.

B.2. Reconstruction evaluation

In the AD model, there's no shape encoder to produce latent variables so we add an in-loop training process to optimize shape latent variables, where we freeze the decoder parameters and optimize latent variables for each test shape. In the VAE training, we also add some refinement steps on the latent variable optimization where we freeze the decoder. We apply this refinement step to both methods w/w.o ARAPReg.

C. More Results

In this section, we show more results of reconstruction (Fig.9), interpolation (Fig.10) and extrapolation (Fig.11) of our methods in variational auto-encoder (VAE) and auto-decoder (AD) frameworks, with and without ARAPReg. We also show more closest shapes for randomly generated shapes in VAE framework with ARAPReg in Fig. 12.

D. Proofs of Propositions in Section 4.2

D.1. Proof of Prop.1

For a shape $\mathbf{g} \in \mathcal{R}^{3n}$ with an infinitesimal vertex displacement $\mathbf{x} \in \mathcal{R}^{3n}$ and $\|\mathbf{x}\|_2 \leq \epsilon$, the local rigidity energy is

$$E(\mathbf{g}, \mathbf{x}) = \min_{\{A_i \in SO(3)\}} \sum_{(i,j) \in \mathcal{E}} w_{ij} \|(A_i - I_3)(\mathbf{g}_i - \mathbf{g}_j) - (\mathbf{x}_i - \mathbf{x}_j)\|^2 \quad (14)$$

where A_i is a 3D rotation matrix denoting the local rotation from $\mathbf{g}_i - \mathbf{g}_j$ to $(\mathbf{g}_i + \mathbf{x}_i) - (\mathbf{g}_j + \mathbf{x}_j)$. Note that here vector indexing is vertex indexing, where $\mathbf{g}_i = \mathbf{g}_{3i:3(i+1)}$.

Since the zero and first-order derivatives from E to $\mathbf{x}$ around zero is 0:

$$E(\mathbf{g}, \mathbf{x})|_{\mathbf{x}=\mathbf{0}} = 0, \quad \frac{\partial E(\mathbf{g}, \mathbf{x})}{\partial \mathbf{x}}|_{\mathbf{x}=\mathbf{0}} = \mathbf{0} \quad (15)$$

We can use second-order Taylor expansion to approximate the energy E when $\mathbf{x}$ is around zero:

$$E(\mathbf{g}, \mathbf{x}) \approx \frac{1}{2} \mathbf{x}^T \frac{\partial^2 E}{\partial \mathbf{x}^2} \mathbf{x} \quad (16)$$

Proposition 3 Given a function $g(\mathbf{x}) = \min_{\mathbf{y}} f(\mathbf{x}, \mathbf{y})$, and define $\mathbf{y}(\mathbf{x}) = (\argmin)_{\mathbf{y}} f(\mathbf{x}, \mathbf{y})$ such that $g(\mathbf{x}) = f(\mathbf{x}, \mathbf{y}(\mathbf{x}))$,

$$\frac{\partial^2 g}{\partial \mathbf{x}^2} = \frac{\partial^2 f}{\partial \mathbf{x}^2} - \frac{\partial^2 f}{\partial \mathbf{x} \partial \mathbf{y}} \left(\frac{\partial^2 f}{\partial \mathbf{y}^2} \right)^{-1} \frac{\partial^2 f}{\partial \mathbf{y} \partial \mathbf{x}} \quad (17)$$

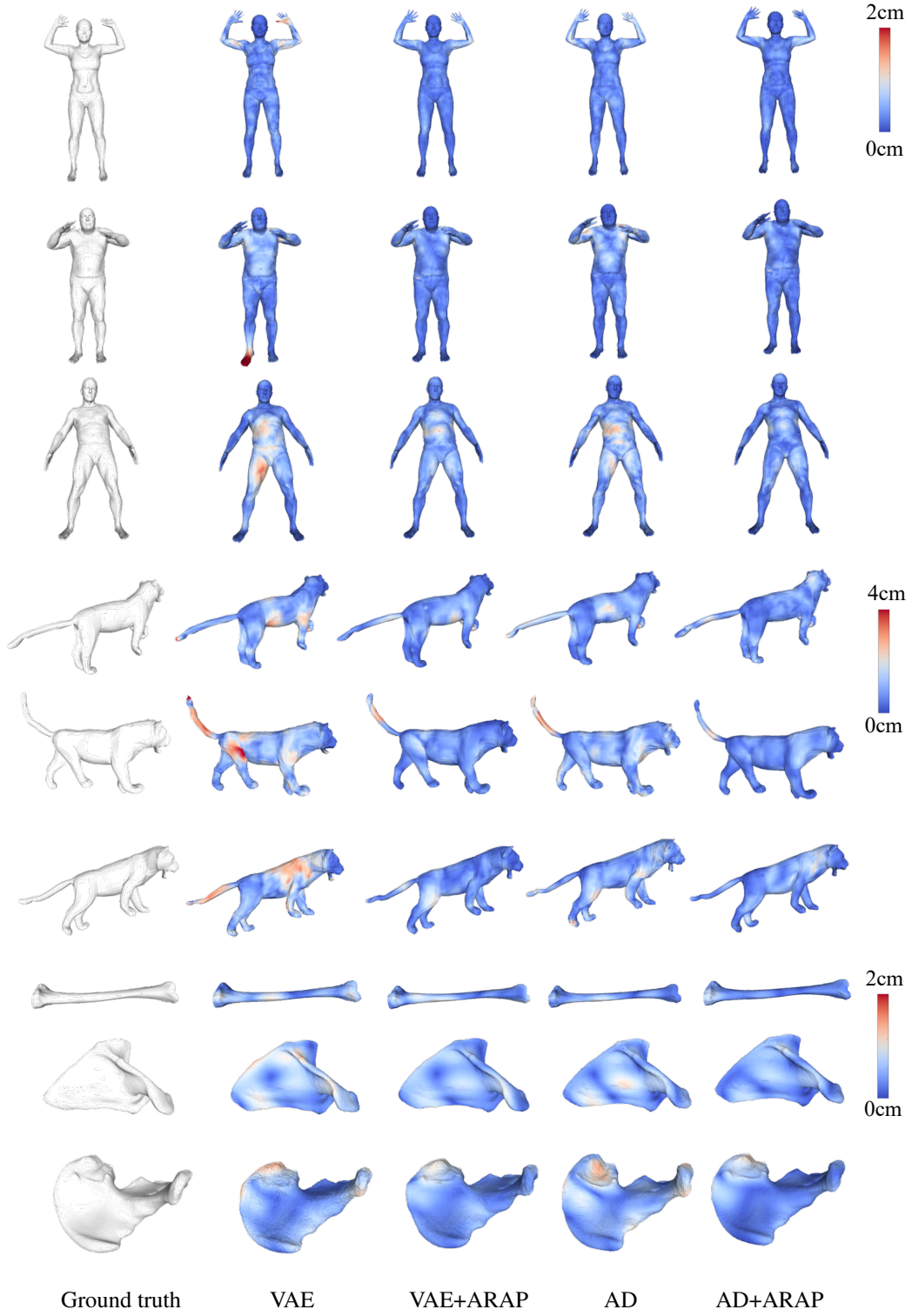


Figure 9: More qualitative results of reconstruction. We show results using VAE and AD generator w/o ARAPReg.

By treating each A_i as a function of $\mathbf{x}$, we can rewrite our energy as

$$E(\mathbf{g}, \mathbf{x}) = f_{\mathbf{g}}(\mathbf{x}, A(\mathbf{x})) \quad (18)$$

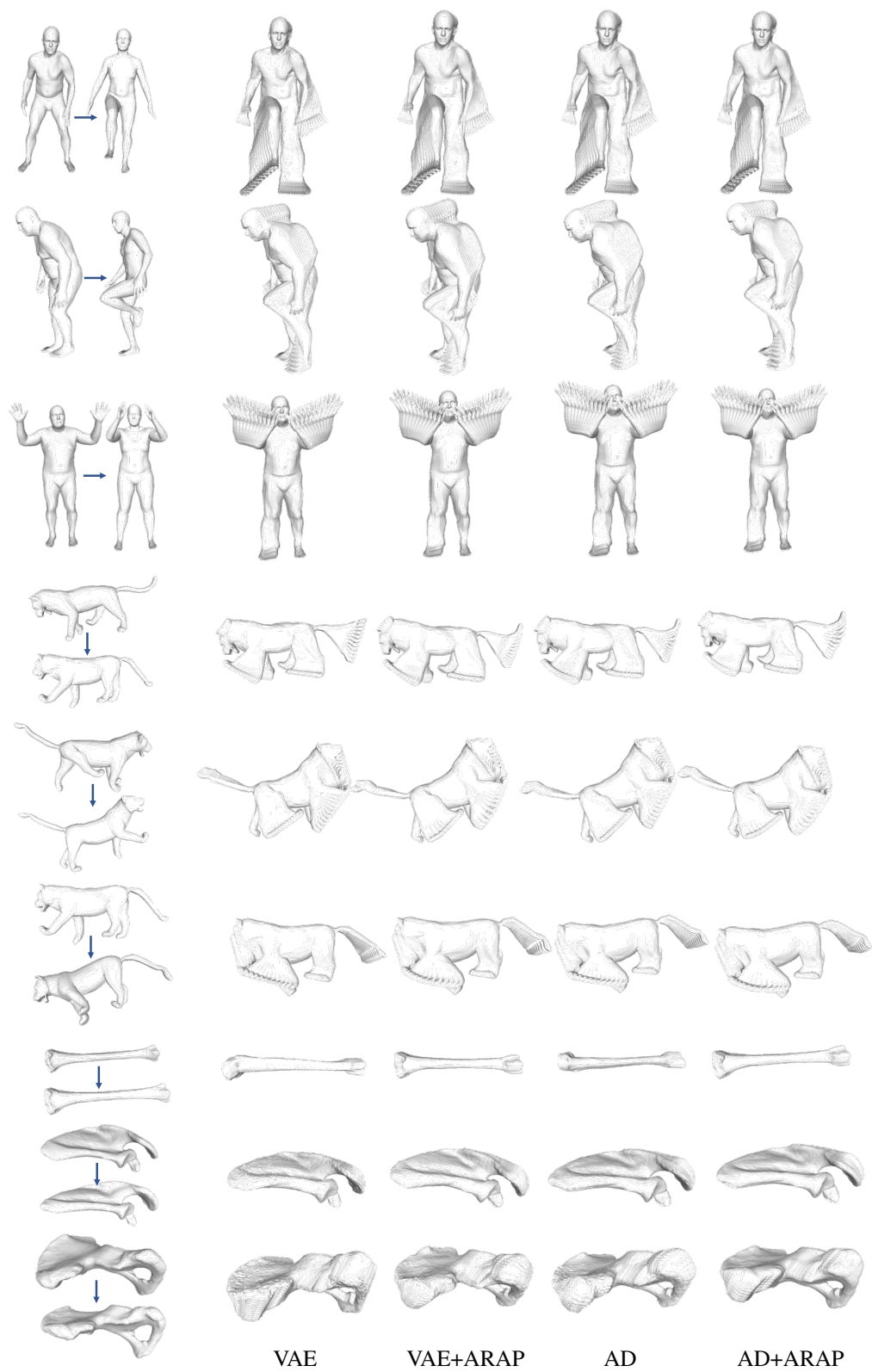


Figure 10: More interpolation results. We show results using VAE and AD generator w/w.o ARAPReg.

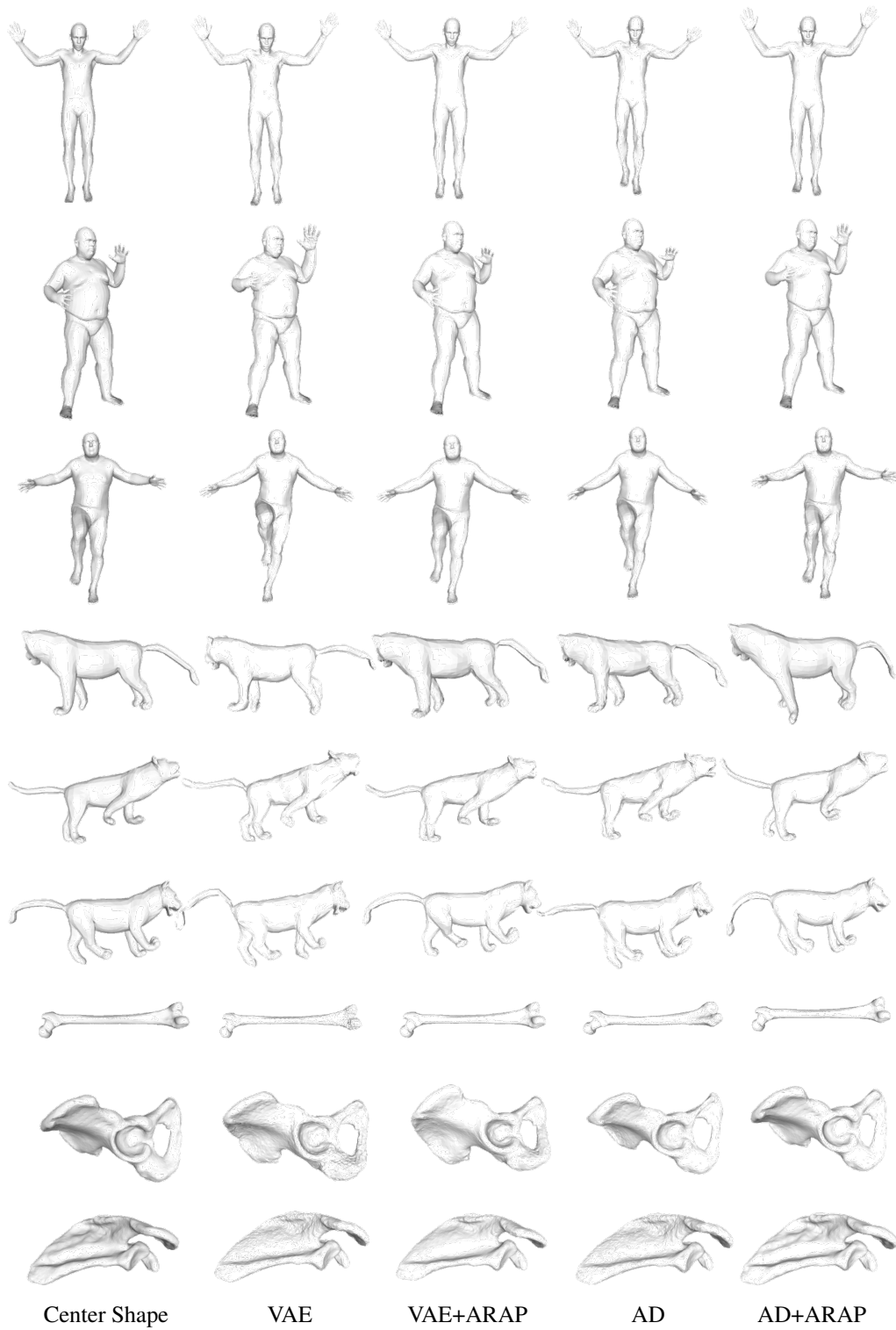


Figure 11: More extrapolation results. We show results using VAE and AD generator w/w.o ARAPReg.

where A is the collection of all A_i .

By using Prop.3, we can get the Hessian from E to $\mathbf{x}$.

In the above formulation, A_i is in the implicit form of $\mathbf{x}$. Now we use Rodrigues' rotation formula to write is explic-

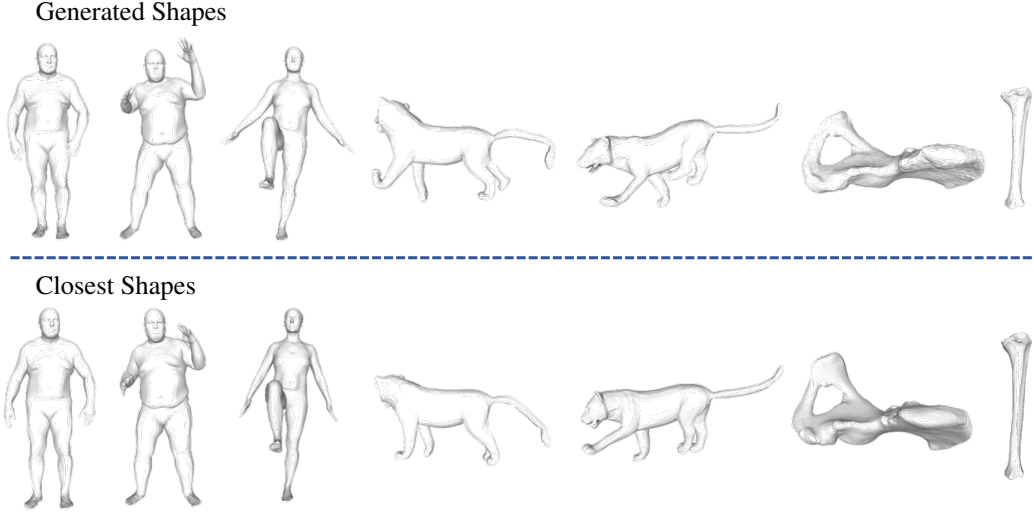


Figure 12: Randomly generated shapes from our VAE framework and their closest shapes in the training set.

itly. For a rotation around an unit axis k with an angle θ , its rotation matrix is

$$A_i = I + \sin\theta \mathbf{k} \times + (1 - \cos\theta)(\mathbf{k} \times)^2 \quad (19)$$

where $\mathbf{k} \times$ is the cross product matrix of vector $\mathbf{k}$.

Since here we apply infinitesimal vertex displacement, rotation angle θ is also infinitesimal. We can approximate 19 as

$$A_i \approx I + \theta \mathbf{k} \times + \frac{1}{2}(\theta \mathbf{k} \times)^2 \quad (20)$$

Let $\mathbf{c} = \theta \mathbf{k}$ and only preserve the first two terms:

$$E(\mathbf{x}) \approx \min_{\{\mathbf{c}_i\}} \sum_{(i,j) \in \mathcal{E}} w_{ij} \|\mathbf{c}_i \times \mathbf{e}_{ij} - (\mathbf{x}_i - \mathbf{x}_j)\|^2 \quad (21)$$

$$= \min_{\{\mathbf{c}_i\}} \sum_{(i,j) \in \mathcal{E}} w_{ij} \|\mathbf{e}_{ij} \times \mathbf{c}_i + (\mathbf{x}_i - \mathbf{x}_j)\|^2 \quad (22)$$

where $\mathbf{e}_{ij} = \mathbf{p}_i - \mathbf{p}_j$

From Prop. 3, we can compute the hessian from E to $\mathbf{x}$ by writing $E(\mathbf{g}, \mathbf{x}) = f_{\mathbf{g}}(\mathbf{x}, \mathbf{c}(\mathbf{x}))$.

We rewrite our energy function in matrix form

$$E = [\mathbf{x}^T \quad \mathbf{c}^T] \begin{pmatrix} L \otimes I_3 & B \\ B^T & C \end{pmatrix} \begin{bmatrix} \mathbf{x} \\ \mathbf{c} \end{bmatrix} \quad (23)$$

where $\otimes$ denotes the kronecker product or tensor product.

The Hessian from E to $\mathbf{x}$ around zero is

$$H_R(\mathbf{g}) = L \otimes I_3 - B^T C^{-1} B \quad (24)$$

Now we compute each term of $H_R(\mathbf{g})$. Expand

$f_{\mathbf{g}}(\mathbf{x}, \mathbf{c}(\mathbf{x}))$:

$$\begin{aligned} f(\mathbf{x}, \mathbf{c}(\mathbf{x})) &= \sum_{(i,j) \in \mathcal{E}} w_{ij} \|\mathbf{e}_{ij} \times \mathbf{c}_i + (\mathbf{x}_i - \mathbf{x}_j)\|^2 \\ &= \sum_{(i,j) \in \mathcal{E}} w_{ij} (\mathbf{x}_i^2 + \mathbf{x}_j^2 - 2\mathbf{x}_i \mathbf{x}_j + 2(\mathbf{e}_{ij} \times \mathbf{c}_i)^T (\mathbf{x}_i - \mathbf{x}_j) \\ &\quad + (\mathbf{e}_{ij} \times \mathbf{c}_i)^T (\mathbf{e}_{ij} \times \mathbf{c}_i)) \end{aligned}$$

L is the weighted graph Laplacian,

$$\mathbf{L}_{ij} = \begin{cases} \sum_{k \in \mathcal{N}_i} w_{ik}, & i = j \\ -w_{ij}, & i \neq j \text{ and } (i, j) \in \mathcal{E} \\ 0, & \text{otherwise} \end{cases} \quad (25)$$

The matrix $\mathbf{B}$ is a block matrix whose 3×3 blocks are defined as

$$B_{ij} = \begin{cases} \sum_{k \in \mathcal{N}_i} w_{ik} \mathbf{e}_{ik} \times, & i = j \\ -w_{ij} \mathbf{e}_{ij} \times, & i \neq j, (i, j) \in \mathcal{E} \\ 0, & \text{otherwise} \end{cases} \quad (26)$$

Finally, $\mathbf{C} = \text{diag}(C_1 \dots C_{|P|})$ is a block diagonal matrix

$$C_i = \sum_{j \in \mathcal{N}_i} w_{ij} (\mathbf{e}_{ij} \times)^T (\mathbf{e}_{ij} \times) \quad (27)$$

$$= \sum_{j \in \mathcal{N}_i} w_{ij} \|\mathbf{e}_{ij}\|_2^2 \mathbf{I}_3 - \mathbf{e}_{ij} \mathbf{e}_{ij}^T \quad (28)$$

which ends the proof. $\square$

D.2. Proof of Prop.2

Consider the eigen-decomposition of

$$\overline{H}_R(\mathbf{g}, J) := U \Lambda U^T,$$

where

$$\Lambda = \text{diag}(\lambda_1(\bar{H}_R(\mathbf{g}, J)), \dots, \lambda_k(\bar{H}_R(\mathbf{g}, J))).$$

Let $\bar{\mathbf{y}} = U^T \mathbf{y}$. Then

$$\begin{aligned} \int_{\mathbf{y}} \mathbf{y}^T \bar{H}_R(\mathbf{g}, J) \mathbf{y} &= \int_{\bar{\mathbf{y}}} \bar{\mathbf{y}}^T \Lambda \bar{\mathbf{y}} = \int_{\bar{\mathbf{y}}} \sum_{i=1}^k \lambda_i(\bar{H}_R(\mathbf{g}, J)) \bar{y}_i^2 \\ &= \sum_{i=1}^k \lambda_i(\bar{H}_R(\mathbf{g}, J)) \int_{\bar{\mathbf{y}}} \bar{y}_i^2 d\bar{\mathbf{y}} \\ &= \frac{1}{k} \sum_{i=1}^k \lambda_i(\bar{H}_R(\mathbf{g}, J)) \int_{\bar{\mathbf{y}}} \sum_{i=1}^k \bar{y}_i^2 d\bar{\mathbf{y}} \\ &= \frac{\text{Vol}(S^k)}{k} \sum_{i=1}^k \lambda_i(\bar{H}_R(\mathbf{g}, J)). \end{aligned}$$

□

E. Gradient of Loss Terms

This section presents the gradients of the loss to the rigidity term.

For simplicity, we will express formulas for gradient computation using differentials. Moreover, we will again replace $\mathbf{g}^\theta$ and $\frac{\partial \mathbf{g}^\theta}{\partial \mathbf{z}}(\mathbf{z})$ with $\mathbf{g}$ and J whenever it is possible. The following proposition relates the differential of $r_R(\mathbf{g}, J)$ with that of $\bar{H}_R(\mathbf{g}, J)$.

Proposition 4

$$dr_R(\mathbf{g}, J) = \alpha \sum_{i=1}^k \frac{\mathbf{u}_i^T d(\bar{H}_R(\mathbf{g}, J)) \mathbf{u}_i}{\lambda_i^{1-\alpha}(\bar{H}_R(\mathbf{g}, J))}. \quad (29)$$

Recall that λ_i and $\mathbf{u}_i$ are eigenvalues of eigenvectors of $\bar{H}_R(\mathbf{g}, J)$.

Proof: The proof is straight-forward using the gradient of the eigenvalues of a matrix, i.e.,

$$d\lambda = \mathbf{u}^T dH \mathbf{u}$$

where $\mathbf{u}$ is the eigenvector of H with eigenvalue λ . The rest of the proof follows from the chain rule. □

We proceed to describe the explicit formula for computing the derivatives of $\mathbf{u}_i^T d(\bar{H}_R(\mathbf{g}, J)) \mathbf{u}_i$. First of all, applying the chain rule leads to

$$\begin{aligned} \mathbf{u}_i^T d(\bar{H}_R(\mathbf{g}, J)) \mathbf{u}_i &= 2 \left((J\mathbf{u}_i)^T H_R(\mathbf{g}) (dJ \cdot \mathbf{u}_i) \right. \\ &\quad \left. - (A(\mathbf{g})J\mathbf{u}_i)^T D(\mathbf{g})^{-1} \cdot (dA(\mathbf{g}) \cdot (J\mathbf{u}_i)) \right) \\ &\quad + (D(\mathbf{g})^{-1} A(\mathbf{g})J\mathbf{u}_i)^T dD(\mathbf{g}) (D(\mathbf{g})^{-1} A(\mathbf{g})J\mathbf{u}_i). \end{aligned}$$

It remains to develop formulas for computing $dJ \cdot \mathbf{u}_i$, $dA(\mathbf{g}) \cdot (J\mathbf{u}_i)$, and $dD(\mathbf{g})$. Note that $J = \frac{\partial \mathbf{g}^\theta}{\partial \mathbf{z}}(\mathbf{z})$. We

use numerical gradients to compute $dJ \cdot \mathbf{u}_i$, which avoid computing costly second derivatives of the generator:

$$d\left(\frac{\partial \mathbf{g}^\theta}{\partial \mathbf{z}}(\mathbf{z})\right) \cdot \mathbf{u}_i \approx \sum_{l=1}^k u_{il} (d\mathbf{g}^\theta(\mathbf{z} + s\mathbf{e}_l) - d\mathbf{g}^\theta(\mathbf{z})) \quad (30)$$

where $s = 0.05$ is the same hyper-parameter used in defining the generator smoothness term; $\mathbf{e}_l$ is the l -th canonical basis of $\mathcal{R}^k$; u_{il} is the l -th element of $\mathbf{u}_i$.

The following proposition provides the formulas for computing the derivatives that involve $A(\mathbf{g})$ and $D(\mathbf{g})$.

Proposition 5

$$dA(\mathbf{g}) \cdot (J\mathbf{u}_i) = -A(J\mathbf{u}_i) \cdot d\mathbf{g}$$

$$\begin{aligned} \mathbf{c}^T dD(\mathbf{g}) \cdot \mathbf{c} &= 2 \sum_{i=1}^n \sum_{k \in \mathcal{N}(i)} \left((\mathbf{g}_i - \mathbf{g}_k)^T (d\mathbf{g}_i - d\mathbf{g}_k) \|\mathbf{c}_i\|^2 \right. \\ &\quad \left. - (\mathbf{c}_i^T (d\mathbf{g}_i - d\mathbf{g}_k)) \cdot ((\mathbf{g}_i - \mathbf{g}_k)^T \mathbf{c}_i) \right) \end{aligned} \quad (31)$$

Proof:

(1). $dA(\mathbf{g}) \cdot (J\mathbf{u}_i)$:

Let's denote $J\mathbf{u}_i$ as $\mathbf{a}$. Now we prove $(A(\mathbf{g}) \cdot \mathbf{a}) = (A(\mathbf{a}) \cdot \mathbf{g})$. Then we will have $d(A(\mathbf{g}))vJ\mathbf{u}_i = d(A(\mathbf{g}) \cdot J\mathbf{u}_i) = d(A(J\mathbf{u}_i) \cdot \mathbf{g}) = A(J\mathbf{u}_i) \cdot d(\mathbf{g})$.

$$\begin{aligned} (A(\mathbf{g})\mathbf{a})_i &= \sum_j A_{ij}(\mathbf{g})\mathbf{a}_j \\ &= \sum_{k \in N(i)} \mathbf{v}_i k \times (\mathbf{a}_i - \mathbf{a}_k) = \sum_{k \in N(i)} \mathbf{v}_i k \times \mathbf{a}_i k \\ &= - \sum_{k \in N(i)} \mathbf{a}_i k \times \mathbf{v}_i k = \sum_j A_{ij}(\mathbf{a})\mathbf{g}_j \\ &= (A(\mathbf{a})\mathbf{g})_i \end{aligned}$$

This finishes the proof.

(2). $\mathbf{c}^T dD(\mathbf{g}) \cdot \mathbf{c}$:

We have $\mathbf{c}_i^T D_{ii}(\mathbf{g}) \cdot \mathbf{c}_i = \sum_{k \in N(i)} (\|\mathbf{v}_{ik}\|^2 \|\mathbf{c}_i\|^2 - \mathbf{c}_i^T \mathbf{v}_{ik} \mathbf{v}_{ik}^T \cdot \mathbf{c}_i)$. We only need to compute the gradient of $\|\mathbf{v}_{ik}\|^2$ and $\mathbf{v}_{ik} \mathbf{v}_{ik}^T$. Note that $\|\mathbf{v}_{ik}\|^2 = \mathbf{v}_{ik}^T \mathbf{v}_{ik}$.

For a vector $\mathbf{a}$, we have $d(\mathbf{a}^T \mathbf{a}) = d(\mathbf{a}^T) \mathbf{a} + \mathbf{a}^T d(\mathbf{a}) = d(\mathbf{a})^T \mathbf{a} + \mathbf{a}^T d(\mathbf{a}) = 2\mathbf{a}^T d(\mathbf{a})$ and similarly, $d(\mathbf{a} \mathbf{a}^T) = 2d(\mathbf{a}) \mathbf{a}^T$. We use these two results to our derivation and we will get the results above.

$$\begin{aligned} \mathbf{c}^T dD(\mathbf{g}) \cdot \mathbf{c} &= \sum_i \sum_{k \in N(i)} (d(\|\mathbf{v}_{ik}\|^2) \|\mathbf{c}_i\|^2 - \mathbf{c}_i^T d(\mathbf{v}_{ik} \mathbf{v}_{ik}^T) \cdot \mathbf{c}_i) \\ &= \sum_i \sum_{k \in N(i)} (d(\mathbf{v}_{ik}^T \mathbf{v}_{ik}) \|\mathbf{c}_i\|^2 - \mathbf{c}_i^T d(\mathbf{v}_{ik} \mathbf{v}_{ik}^T) \cdot \mathbf{c}_i) \\ &= \sum_i \sum_{k \in N(i)} 2 \left((\mathbf{g}_i - \mathbf{g}_k)^T (d\mathbf{g}_i - d\mathbf{g}_k) \|\mathbf{c}_i\|^2 \right. \\ &\quad \left. - (\mathbf{c}_i^T (d\mathbf{g}_i - d\mathbf{g}_k)) \cdot ((\mathbf{g}_i - \mathbf{g}_k)^T \mathbf{c}_i) \right) \end{aligned}$$