



Multi-directional geodesic neural networks via equivariant convolution

Adrien Poulenard, Maks Ovsjanikov

► To cite this version:

Adrien Poulenard, Maks Ovsjanikov. Multi-directional geodesic neural networks via equivariant convolution. ACM Transactions on Graphics, 2018, 37 (6), pp.236. 10.1145/3272127.3275102 . hal-04479680

HAL Id: hal-04479680

<https://hal.science/hal-04479680>

Submitted on 27 Feb 2024

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.

Multi-directional Geodesic Neural Networks via Equivariant Convolution

ADRIEN POULENARD, LIX, Ecole Polytechnique
MAKS OVSJANIKOV, LIX, Ecole Polytechnique

We propose a novel approach for performing convolution of signals on curved surfaces and show its utility in a variety of geometric deep learning applications. Key to our construction is the notion of directional functions defined on the surface, which extend the classic real-valued signals and which can be naturally convolved with real-valued template functions. As a result, rather than trying to fix a canonical orientation or only keeping the maximal response across all alignments of a 2D template at every point of the surface, as done in previous works, we show how information across all rotations can be kept across different layers of the neural network. Our construction, which we call *multi-directional geodesic convolution*, or *directional convolution* for short, allows, in particular, to propagate and relate directional information across layers and thus different regions on the shape. We first define directional convolution in the continuous setting, prove its key properties and then show how it can be implemented in practice, for shapes represented as triangle meshes. We evaluate directional convolution in a wide variety of learning scenarios ranging from classification of signals on surfaces, to shape segmentation and shape matching, where we show a significant improvement over several baselines.

CCS Concepts: • **Computing methodologies** → **Shape analysis**;

Additional Key Words and Phrases: Geometric Deep Learning, Convolution, Rotation Equivariance, Parallel Transport

ACM Reference Format:

Adrien Poulenard and Maks Ovsjanikov. 2018. Multi-directional Geodesic Neural Networks via Equivariant Convolution. *ACM Trans. Graph.* 37, 6, Article 236 (November 2018), 14 pages. <https://doi.org/10.1145/3272127.3275102>

1 INTRODUCTION

The success of convolutional neural networks (CNNs) for image processing tasks [Krizhevsky et al. 2012] has brought attention from the geometry processing community. In recent years multiple techniques have been developed to reproduce the success of CNNs in the context of geometry of curved surfaces with applications including shape recognition [Su et al. 2015], segmentation [Kalogerakis et al. 2017; Maron et al. 2017] or shape matching [Boscaini et al. 2015], among many others. A key aspect of CNNs is that they rely on convolution operations. In the Euclidean domain the notion of convolution is well-defined whereas on non-Euclidean spaces,

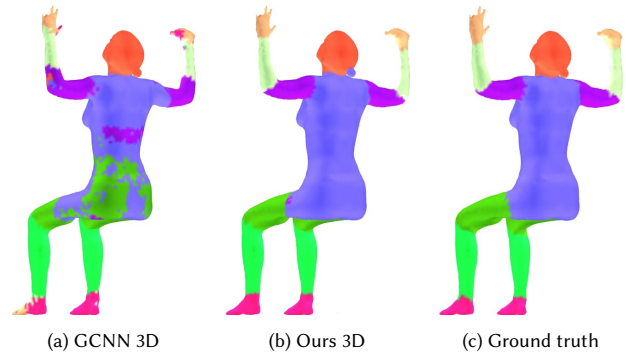


Fig. 1. Our approach (middle) enables a variety of learning applications, including shape segmentation using “weak” signals, such as 3D coordinates as input, via a richer and more accurate notion of convolution on surfaces. This leads to practical improvements compared to previous methods such as GCNN [Boscaini et al. 2015] as shown on the left.

in general, there is no direct analogue that satisfies all the same properties.

Different approaches have been proposed to overcome this limitation. Perhaps the simplest and most common consist in either performing convolution directly on the surrounding Euclidean 3D space (using so-called volumetric approaches [Maturana and Scherer 2015], [Qi et al. 2016]), or constructing multiple projections (views) of an embedded object from different angles and applying standard CNNs in 2D [Su et al. 2015]. This idea has also been extended to using more general mappings onto canonical domains, including the plane [Ezuz et al. 2017; Sinha et al. 2016], or e.g. toric domains which admit a global parameterization and where convolution can be defined naturally [Maron et al. 2017]. Unfortunately, such mappings can induce significant distortion and might be restricted to only certain topological classes.

On the other hand, several *intrinsic* approaches have been proposed to define analogues of convolution directly on the manifold [Boscaini et al. 2015, 2016; Masci et al. 2015]. These techniques aim at learning local template (or kernel) functions, which can be mapped onto a neighborhood of each point on a surface and convolved with signals defined on the shape. These methods are general, can be applied regardless of the shape topology and moreover are not sensitive to the changes in shape embedding, making them attractive in non-rigid shape matching applications, for example. Unfortunately, general surfaces lack even *local* canonical coordinate systems, which means that the mapping of a template onto the surface is defined only up to the choice of an orthonormal basis of the tangent plane at every point. To overcome this limitation, previous methods either only consider the maximal response over a certain number of rotations [Boscaini et al. 2015] or aim to resolve these ambiguities

Authors' addresses: Adrien Poulenard, LIX, Ecole Polytechnique, adrien.poulenard@inria.fr; Maks Ovsjanikov, LIX, Ecole Polytechnique, maks@lix.polytechnique.fr.

using principal curvature directions [Boscaini et al. 2016; Masci et al. 2015]. Unfortunately such approaches can either lead to more instabilities or, in the case of angular max pooling, lose the relative orientation of the kernels across multiple convolutional layers.

In this paper, we propose to overcome this key limitation of intrinsic methods by aligning the convolutional layers of the network. Our idea is to consider *directional* (or, equivalently, *angular*) functions defined on surfaces, and to define a notion of convolution for them which results, again, in *directional functions* without loss of information. This allows us to impose specific canonical relations across the layers of a neural network, lifting the directional ambiguity to only the last layer. We then need to take the maximal response over all rotations *only on the last layer*. This allows us to better capture the relative response at different points, leading to an overall improvement in training and test accuracy.

2 RELATED WORK

Geometric Deep Learning is an emerging field aimed at applying machine learning techniques in the context of geometric data analysis and processing. Below we review the techniques most closely related to ours, and especially concentrate on various ways to define and use convolution on geometric (3D) shapes and refer the interested reader to several recent surveys, e.g. [Bronstein et al. 2017; Xu et al. 2016].

2.1 Extrinsic and Volumetric Techniques

Perhaps the most common approach for exploiting the power of Convolutional Neural Networks (CNNs) on 3D shapes is to transform them into 2D images, for example by rendering multiple views of the object. Some of the earliest variants of this idea include methods that represent shapes as unordered collections of views (or range images) [Kalogerakis et al. 2017; Su et al. 2015; Wei et al. 2016] or exploit the panorama image representation [Sfikas et al. 2017; Shi et al. 2015] among others, as well as techniques based on Geometry Images [Sinha et al. 2016], which represent the 3D geometry by mapping the coordinates onto the plane.

Another common set of methods considers 3D shapes as volumetric objects and defines convolution simply in the Euclidean 3D space, e.g. [Maturana and Scherer 2015; Wu et al. 2015] among others. Due to the potential memory complexity of this approach, several efficient extensions have been proposed, e.g., [Klokov and Lempitsky 2017; Wang et al. 2017], and a comparison between view-based and volumetric approaches has been presented in [Qi et al. 2016].

Other recent techniques analyze 3D shapes simply as collections of points and define deep neural network architectures on point clouds, including, most prominently PointNet [Qi et al. 2017], PointNet++ [Qi et al. 2017] and several extensions such as PointCNN [Li et al. 2018] and Dynamic Graph CNN [Wang et al. 2018] for shape classification and segmentation, and PCPNet [Guerrero et al. 2018] for normal and curvature estimation, among others.

Despite their efficiency and accuracy in certain situations, these techniques rely directly on the embedding of the shapes, and are thus very sensitive to changes in shape pose, which can limit their use, for example in non-rigid shape matching.

2.2 Intrinsic and Graph-based Techniques

To overcome this limitation, several *intrinsic* methods have been proposed, for defining and exploiting convolution directly on the surface of the shape. This includes spectral methods, which exploit the relation between convolution and multiplication in the spectral domain [Boscaini et al. 2015], and which have also been applied on general graphs [Defferrard et al. 2016]. A similar technique, treating shapes as graphs, has been used for analyzing arbitrary shape collections [Yi et al. 2017], while synchronizing their laplacian eigenbases with functional maps [Ovsjanikov et al. 2012]. Another recent method [Ezuz et al. 2017] consists in optimizing an embedding of the shape onto the planar domain using intrinsic metric alignment [Solomon et al. 2016]. Finally, a very recent Surface Networks approach [Kostrikov et al. 2017] is based on stacking layers consisting of combinations of features and their images by the Laplace or Dirac operator to exploit intrinsic and extrinsic information.

More closely related to our approach are techniques based on *local* shape parameterization, which define convolution of a signal with a learned kernel on a region of the shape surface. A seminal work in this direction was done in [Masci et al. 2015] where the authors defined Geodesic Convolutional Neural Networks (GCNNs), which locally align a given kernel with the shape surface at each point, and perform convolution in the tangent plane. Unfortunately, the absence of canonical coordinate systems on surfaces leads to a one-directional ambiguity in the alignment. To rectify this, the authors of [Masci et al. 2015] proposed to take the maximal response across all possible alignments. Several later extensions of this approach have used different local patch parameterizations, [Boscaini et al. 2016; Monti et al. 2017] and also used principal curvature directions to resolve the directional ambiguity. Unfortunately, principal curvature directions can be highly unstable, and not uniquely defined even on basic domains such as the sphere and the torus, which can lead to over-fitting in the training. Finally, a recent approach has been proposed for defining convolution via mapping onto a toric domain [Maron et al. 2017], which admits a global parameterization.

2.3 Contribution

In our work, we show that the directional ambiguity that exists when mapping a template (kernel) onto the surface, as done in [Boscaini et al. 2016; Masci et al. 2015; Monti et al. 2017] *can be maintained* across the layers of the deep neural network without relying on a canonical direction choice or only keeping the maximal response across all directions. To achieve this, we first extend real-valued signals to more general *directional functions*. We then show that a directional function can be convolved with a template to produce another directional function, and can thus be stacked in a deep neural network. As a result the directional ambiguity is lifted up to the last layer, where it can be resolved by taking the maximum response only once. We demonstrate through extensive experiments that this leads to overall improvement in accuracy and robustness in a range of applications. Finally, we extend previous approaches such as [Boscaini et al. 2016; Monti et al. 2017] by adding spatial pooling layers through mesh simplification and exploiting residual learning (ResNet) blocks [He et al. 2016] in the architecture.

3 CONVOLUTION OVER MANIFOLDS

Throughout our work, we assume that we are dealing with 3D shapes, represented as oriented (manifold) 2D surfaces, embedded in 3D. For simplicity of the discussion, we also assume that the shapes are without boundary, although our practical implementation does not have this limitation.

Throughout our discussion, the input signal is assumed to be a tuple of real-valued functions defined on each shape in the collection. These can either represent some geometric descriptors, or even simply the 3D coordinates of each point. In this section, we describe how the convolution operation is applied to a given signal for a fixed template. Our approach follows the general structure proposed in [Masci et al. 2015], but we highlight the key differences, arising from our use of *directional convolution*. Finally, let us note that for simplicity we concentrate on oriented two-dimensional manifolds, although most of our discussion can be adapted easily to more general settings.

Recall that in the standard two-dimensional Euclidean setting, the cross-correlation or convolution of a real-valued function f by a smooth compactly supported (template) function $k : \mathbb{R}^2 \rightarrow \mathbb{R}$ is defined for all $x \in \mathbb{R}^2$ by:

$$(f * k)(x) := \int_{\mathbb{R}^2} f(t)k(t+x)dt.$$

Convolution is a way to compare the function f to a template function k at every point. We can reinterpret convolution in the following way: first, for every point x we identify the tangent space to x with a copy of $\mathbb{R}^2$ whose origin has been shifted to x . Then, translation by x can be seen as linear isometry:

$$\tau_x : T_x \mathbb{R}^2 \rightarrow \mathbb{R}^2 \simeq T_0 \mathbb{R}^2, \quad \tau_x(p) = p + x \quad \forall p,$$

where $\mathbb{R}^2 \simeq T_0 \mathbb{R}^2$ simply means that the tangent plane at the origin can be identified with the whole of $\mathbb{R}^2$. Moreover, the identity map “centered at x ” can be seen as another isometry:

$$\text{Id}_x : T_x \mathbb{R}^2 \rightarrow \mathbb{R}^2, \quad \text{Id}_x(p) = p \quad \forall p.$$

Therefore can rewrite the convolution of f by k at x , as:

$$(f * k)(x) = \langle \text{Id}_x^* f, \tau_x^* k \rangle_{L^2},$$

where the superscript $*$ means the pull-back of the function with respect to the map, and $\langle \cdot, \cdot \rangle_{L^2}$ is the standard L^2 inner product. We can now generalize this construction to two-dimensional manifolds by generalizing the maps Id_x^* and τ_x^* . At every point x of a Riemannian manifold X the exponential map:

$$\exp_x^X : T_x X \rightarrow X$$

generalizes the previous map $\text{Id}_x : T_x \mathbb{R}^2 \rightarrow \mathbb{R}^2$ in the sense that $\exp_x^{\mathbb{R}^2} = \text{Id}_x$. Following [Masci et al. 2015], we assume that the template function k is defined over a copy of $\mathbb{R}^2$, denoted by $T_0 \mathbb{R}^2$, and generalize the map τ_x by isometrically aligning this copy with the tangent plane $T_x X$ at x . That is, the map $\tau_x : T_x X \rightarrow T_0 \mathbb{R}^2$ must be a linear isometry. This map is uniquely-defined given a choice of the correspondence of the coordinate axes of $\mathbb{R}^2$ with a coordinate frame of $T_x X$. For oriented surfaces, this reduces to the alignment of one coordinate axis. If $e_1 := (1, 0)$ is the first coordinate axis on $\mathbb{R}^2$, this means that the inverse of the map τ must send e_1 to all tangent spaces, and that k can be pulled-back onto $T_x X$ for any x , via $\tau_x^* k$.

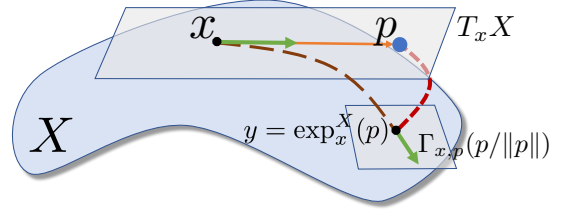


Fig. 2. The *completed* exponential map at a point $x \in X$ sends any non null tangent vector $p \in T_x X$ to the couple $(y = \exp_x^X(p), v = \Gamma_{x,p}(\frac{p}{\|p\|}))$ where $v \in T_y X$ is the result of parallel transport of $\frac{p}{\|p\|}$ along the geodesic from x to y , with the initial direction given by p .

Since unlike the Euclidean case, there is no global choice of reference direction on a surface, an arbitrary choice of the pre-image of e_1 on every tangent space $T_x X$ can lead to biased results, which furthermore will not generalize across different (e.g., training and test) shapes.

To resolve this ambiguity, the authors of [Masci et al. 2015] consider the family of maps $\tau_{x,v}$ parameterized by a choice of a unit vector v in the tangent plane of x , which is mapped to the first coordinate axis in $\mathbb{R}^2$, i.e. $\tau_{x,v}(v) = e_1$. They then define geodesic convolution by taking the maximum response of the signal to the template function k mapped via $\tau_{x,v}$ across all choices of v .

$$(f \circledast k)(x) = \max_v \langle (\exp_x^X)^* f, \tau_{x,v}^* k \rangle_{L^2}. \quad (1)$$

The two main advantages of this procedure are that: 1) it does not depend on the choice of reference direction in the tangent planes, and 2) the output of $f \circledast k$ is again a real-valued function, so that convolutions can be applied successively within a deep network.

Unfortunately, only keeping the maximum response also results in a loss of directional information, which can make it more difficult to detect certain types of features in the signal. In particular, since the maximum is applied independently at every point, directional information of the response of the signal to the template is not shared across nearby points.

3.1 Convolution of Directional Functions

We propose to address these limitations by considering convolution of a template function with a more general notion of *directional functions* defined on arbitrary surfaces.

We call a directional function, any function $\varphi(x, v)$ that depends on both the point x on a surface X , and on the unit direction v in the tangent plane $T_x X$ at x . Clearly any real-valued function $f : X \rightarrow \mathbb{R}$ can be lifted to a directional function $\tilde{f}$, simply by ignoring the directional argument and setting $\tilde{f}(x, v) = f(x)$ for any v .

Our key observation is that the convolution of a *directional function* φ with respect to a template $k : T_0 \mathbb{R}^2 \rightarrow \mathbb{R}$ can be defined naturally, so that the result of the convolution is, once again, a directional function. For this we first complete the exponential map, so that for any point $p \in T_x X$ in the tangent plane of x , it produces both a point y on X and a unit direction in the tangent plane of y . To achieve this we define the completed exponential map:

$$\overline{\exp}_x^X(p) = (\exp_x^X(p), \Gamma_{x,p}(p/\|p\|)), \quad (2)$$

where $\Gamma_{x,p}(p/\|p\|)$ is the parallel transport of the unit vector $p/\|p\|$ from the tangent plane of x to the tangent plane of $y = \exp_x^X(p)$ along the geodesic between x and y with initial velocity $p/\|p\|$ (see Figure 2). Thus, for any point p in the tangent plane of x , $\overline{\exp}_x^X(p)$ outputs a point y on the manifold and a vector in the tangent plane of y . Moreover, since parallel transport along geodesics preserves the norms of vectors, $\Gamma_{x,p}(p/\|p\|)$ must also be a unit vector. This map is well-defined everywhere, except at the origin $p = 0$. This does not pose a problem in our setting, however, since we only use this map inside an integral. We now define *multi-directional geodesic convolution*, or *directional convolution* for short, of a template $k : T_0\mathbb{R}^2 \rightarrow \mathbb{R}$ with a directional function φ by:

$$(\varphi \star k)(x, v) = \langle (\overline{\exp}_x^X)^* \varphi, \tau_{x,v}^* k \rangle_{L^2}. \quad (3)$$

Note that, $\psi_x = (\overline{\exp}_x^X)^* \varphi$ is a real-valued function in the tangent plane of x , where $\psi_x(p) = \varphi(\overline{\exp}_x^X(p))$ (see Figure 5). Moreover, note that the result of a convolution of a directional function φ and a template k is once again a directional function, as it depends both on the point x and on the direction v . We also remark that directional convolution can be extended to regular-valued functions. Thus, for any $f : X \rightarrow \mathbb{R}$ we simply set: $f \star k := \hat{f} \star k$ where $\hat{f}$ is the lifting of f to a directional function, as described above.

3.2 Directional vs. Geodesic Convolution

As suggested in the previous section, geodesic convolution introduced in [Masci et al. 2015] and our directional convolution are closely related. Indeed, below we show that directional convolution is strictly more informative than geodesic convolution. The following proposition (proved in the Appendix) shows that geodesic convolution can be factorized by taking the maximal directional response of directional convolution thus losing the directional information.

PROPOSITION 3.1. *Let f a function on X and k a template. Denote by $\tilde{f} = (x, v) \mapsto f(x)$ the directional function obtained via $\tilde{f}(x, v) = f(x)$ for all unitary $v \in T_x X$. Then:*

$$f \otimes k(x) = \max_{v \in T_x X} (\tilde{f} \star k)(x, v)$$

Intuitively, applying directional convolution allows to keep track of the direction that the signal comes from. To illustrate this, we consider the directional convolution of the indicator function 1_x of a point x , by a shifted Dirac kernel $\delta_{(t,0)}$ for some $t > 0$. It is easy to see that the result of $1_{(x,v)} \star \delta_{(t,0)}$ is the indicator function of the set of couples (y, v) such that $y \in X$ is at (geodesic) distance t from x and $v \in T_y X$ points in the direction of x (see Figure 3 for an illustration). Moreover, as also illustrated in Figure 3, applying directional convolution by $\delta_{(t,0)}$ multiple times will propagate the signal along geodesics from the source point x while maintaining the directional information attached to it.

In contrast, after angular max-pooling directional information to the source point is lost. This means that when applying the basic geodesic convolution repeatedly (e.g. when stacking multiple convolution layers in a neural network) the signal does not necessarily

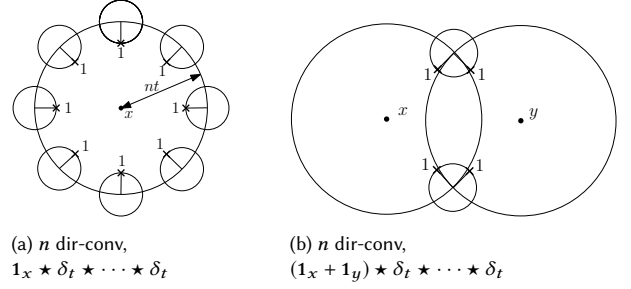


Fig. 3. (left) The result of applying n times the directed convolution by the shifted Dirac kernel $\delta_{(t,0)}$ to the indicator function 1_x of point x . We obtain a response along the circle of radius nt in the directions pointing to x . This allows to detect the presence of a feature at a given distance and to point in the direction of that feature. (right) The response of the convolution with two isolated signals at x and y is located along two circles, if they intersect the intersection points store information about the relative angle and distance between the features at x and y .

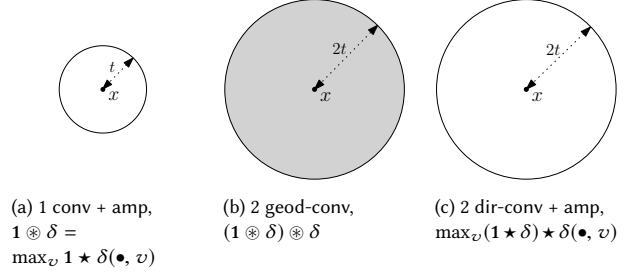


Fig. 4. Comparison between repeated geodesic (geod, middle) and directional convolution (dir, right) of the indicator of a point 1 by a shifted Dirac kernel $\delta = \delta_t$. (b) With directional convolution we obtain a response along the circle of radius $2t$ indicating the presence of a feature at *exactly* distance $2t$ (c) With geodesic convolution, we obtain a response along the full disc of radius $2t$ only indicating the presence of a feature at distance $\leq 2t$.

travel along geodesic paths. Thus, for example, $(1_x \otimes \delta_{(t,0)}) \otimes \delta_{(t,0)}$ is a geodesic ball instead of a circle of radius $2t$ around x (see Figure 4). This might result in a loss of information and makes stacks of filters based on geodesic convolution less efficient in estimating the distance between features or their relative position, compared to directional convolution shown in Figure 3b. We also note briefly that directional geodesic convolution does not admit an identity kernel but identity can be obtained as a limit of convolutions by shifted Dirac kernels since:

$$\lim_{t \rightarrow 0} \overline{\exp}_x^X(tv) = (x, v).$$

3.3 Directional Convolution in Angular Coordinates

The expression in Eq. (3) is coordinate-free, however to implement it we must choose a particular coordinate system to represent directional functions. In practice, we represent tangent vectors in polar coordinates in their tangent plane. This representation implies an arbitrary choice of reference direction $e_x \in T_x X$ in each tangent plane. For any angle $\theta \in [0, 2\pi)$ we denote by $e_x(\theta) = R_\theta e_x$, the

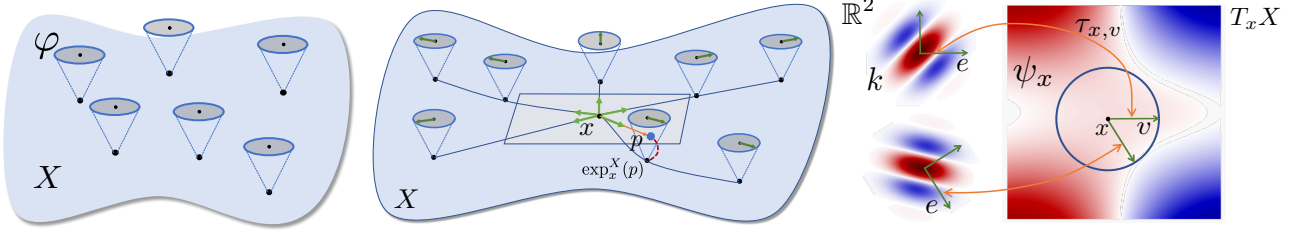


Fig. 5. Overview of our directional convolution: **(left)** we start with a surface X and a directional function φ , which assigns a real value to every unit direction in the tangent plane of every point of X , shown here as a cone above some points. **(middle)** given a point $x \in X$, we pull back φ onto the tangent plane of x by computing for every point $p \in T_x X$, its image $y \in X$ via the exponential map, and the transport of the unit vector $p/\|p\|$ to the tangent plane of $y = \exp_x^X(p)$. We assign to p the value of φ at y and the transported vector. By repeating this for every p we obtain a real-valued function ψ_x over the tangent plane of x . **(right)** Given a unit vector v tangent at x there is a unique orthogonal orientation-preserving linear map $\tau_{x,v}$ from $T_x X$ to $\mathbb{R}^2$ sending v to $e := (1, 0)$. We pull back a kernel k defined over $\mathbb{R}^2$ to $T_x X$ and take its L^2 product with ψ_x giving us a real-number for each unit tangent vector at x . Repeating this procedure for every $x \in X$ we obtain a new directional function which is the directional convolution of φ by the kernel k .

vector obtained by rotating e_x by angle θ in the tangent plane at x . Instead of operating with directional functions, we then work with *angular functions* as follows: given a directional function φ , and a choice of reference direction e_x for each point x , we define an angular function φ_e by:

$$\varphi_e(x, \theta) := \varphi(x, e_x(\theta)).$$

that is, φ_e is just the function φ expressed in polar coordinates in each tangent plane with respect to the reference direction e . Note also that given reference directions $e_x \in T_x X$ for all $x \in X$, both the exponential map and the parallel transport can be expressed for points and vectors described in polar coordinates in each tangent plane (see Appendix). Finally, given the family of reference directions $e = (e_x)_{x \in X}$ and the angular function φ_e we denote by $\varphi_e \star k$ its directional convolution with kernel k w.r.t. e . To define this operation, we first convert φ_e to its directional counterpart, using the reference directions e and then simply apply the definition given in Eq. (3) above and convert it back to an angular function using the same reference directions:

$$\varphi_e \star k := (\varphi \star k)_e$$

The following proposition (proved in the Appendix) shows that the result of directional convolution of angular functions is *equivariant* with respect to the change of the reference directions. Namely:

PROPOSITION 3.2. *Let $e_x \in T_x X$ be a family of unit tangent vectors defined at every point $x \in X$ then for any family of rotations R_x of angle θ_x at point x we have:*

$$(\varphi_{Re} \star k)(x, \theta) = (\varphi \star k)_{Re}(x, \theta) = (\varphi_e \star k)(x, \theta + \theta_x)$$

Where $(Re)_x := R_x e_x$ at each point x .

This proposition guarantees that even if we pick an arbitrary reference direction in the tangent plane of every point x , the result of possibly multiple directional convolution steps is the same up to an angle offset. Moreover, this offset is fixed and is the angle difference between the two reference directions.

When directional convolution operator $\star$ is defined directly in angular coordinates, and angles are discretized in angular bins, as they are in our practical implementation, this implies that changing

the reference direction at a given point by applying a circular permutation to the angular bins' indices will lead to the same permutation of the result of directional convolution. Thus, no further ambiguity is introduced between layers of directional convolution, and the initial rotational ambiguity is lifted to the last layer. We resolve it by applying angular max pooling. This is crucial to learning since it allows learn the same features independently of the reference directions used to define the angular coordinates.

3.4 Multi-Directional Convolutional Neural Networks

Our main motivation for introducing directional convolution is to use it as a layer inside a deep neural network for shape analysis tasks. We define a directional convolution layer simply by replacing the standard convolution (resp. geodesic convolution) operator and regular functions by directional functions in the convolutional layer definition. Crucially, since directional convolution by a template kernel results in another directional function, we can stack *directional convolutional layers* in a neural network.

Since the input of the network is typically given as a regular signal $f : X \rightarrow \mathbb{R}^d$ we first convert it into a directional function, by setting $\tilde{f}(x, v) := f(x)$ for all $x \in X$ and any unit vector $v \in T_x X$. A convolutional neural network has stacks of learnable template kernels $k_{pq}^i : \mathbb{R}^2 \rightarrow \mathbb{R}$, a collection of learnable bias vectors b^i , and activation functions ξ_i . In the simplest setting, we can stack layers sequentially so that the output of $(i + 1)$ st layer of the network is defined by:

$$\begin{cases} y_0(f) = \tilde{f} \\ (y_{i+1})_p := \xi_{i+1}(\sum_q ((y_i)_q \star k_{pq}^i) + b_p^i) \end{cases}$$

Our goal during training is then to learn the parameters k_{pq}^i and b^i so that the output function minimizes some error for a set of examples. In practice, we consider more complex architectures, as described in Section 6. In all cases, angular max pooling is applied to the last layer y_n to obtain a regular signal over X . We call networks based on directional convolution MDGCNN for Multi Directional Geodesic Convolutional Neural Networks.

4 DIRECTIONAL CONVOLUTION OVER DISCRETE SURFACES

We assume that all shapes are represented as connected manifold triangle meshes, possibly with boundary. In this section we describe our general approach for implementing directional convolution in practice. To simplify the presentation, we describe only the main steps necessary for the implementation, and defer the exact implementation details to Section 8.1 in the Appendix.

To implement directional convolution, we need to decide on the discrete representation of template functions, angular functions on the surface, the exponential map and parallel transport. These are described as follows:

4.1 Template Functions

To discretize template functions of the form $k : \mathbb{R}^2 \rightarrow \mathbb{R}$, we represent them via windows of discrete polar coordinates in the plane: (ρ_i, θ_j) , where $\rho_i := \frac{(i+1)R}{N_\rho+1}$ is a set of radii, for i varying between 0 and $N_\rho - 1$ bounded by the window radius R and $\theta_j := \frac{2j\pi}{N_\theta}$ is a uniform discrete sampling of angles $[0, 2\pi)$. This means that each k associates a real value to each pair (ρ_i, θ_j) , and can therefore be stored as a matrix of size N_ρ, N_θ . Note that during training these are the unknowns that need to be trained.

4.2 Angular Functions

Angular functions are real-valued functions that depend on the point on a given surface and on a direction in its tangent plane. For a mesh with N_v vertices, we represent these as matrices of size (N_v, N_θ) , where N_θ defines a discrete angular sampling, in the same way as for template functions above, but with respect to some arbitrary reference direction in the tangent plane of each point.

4.3 Exponential Map

To discretize the exponential map, similarly to previous approaches, we use Geodesic Polar Coordinates (GPC), which associate a window to every point, and represent other points in its neighborhood, through the geodesic distance ρ and angle θ to the base point in its window. We discretize GPC using the same set of discrete polar pairs (ρ_i, θ_j) as for the template functions. This means, in particular, that given a vertex i , the points in its GPC might not fall on the vertices of the mesh. We therefore use barycentric coordinates inside triangles to interpolate values at points in the GPC windows, based on the values at the vertices of the mesh. This procedure is crucial since it helps us gain resilience against changes in mesh structure. We can model the GPC using a tensor E of size $(N_v, N_\rho, N_\theta, N_v)$, such that E_{vijw} stores, the barycentric coordinates with respect to vertex w of the point having polar coordinates ρ_i, θ_j in the GPC of the window of vertex v . Note that by definition, generically, E_{vijw} will have exactly three non-zero values for fixed v, i, j .

4.4 Parallel Transport

Finally, we need to discretize the parallel transport operation in order to define the discrete analogue of the (completed) exponential map $\overline{\text{exp}}$ defined in Eq. (2). Note that a key feature of our definition is that parallel transport needs to be defined only for unit vectors

connecting a point p in the tangent plane of v to the origin, which correspond, in our discretization, to the angular polar coordinate θ of p . This means that we need to compute the angle difference between θ in the window of v and the corresponding angle in the GPC of the exponential map of p . However due to angular discretization the transported angle does not necessarily fall into the angles $(\theta_j)_j$ of the window at p . We therefore interpolate between consecutive discrete angles to compute angular functions. Similarly to the exponential map above, the parallel transport can be discretized as a 5D tensor Γ of size $(N_v, N_\rho, N_\theta, N_v, N_\theta)$ where Γ_{vijwk} stores the interpolating coordinate with respect to the angle θ_k at vertex w of the geodesic parallel transport of angle θ_j to the point with polar coordinates (ρ_i, θ_j) in the GPC of the window of vertex v . The (completed) angular exponential map $\overline{\text{exp}}$ can then be interpolated in both the spatial and angular domains as a 5D tensor $\bar{E}$ of size $(N_v, N_\rho, N_\theta, N_v, N_\theta)$ such that:

$$\bar{E}_{vijwk} := E_{vijw} \Gamma_{vijwk}.$$

With these constructions at hand, discretizing the generalized convolution, defined in Eq. (3) can be done simply by matrix multiplication. Note that aligning the template function with the tangent plane at x via a map τ_x simply corresponds to aligning the angular coordinate of the template with the angular coordinate of the GPC at x . While aligning a rotated template function via $\tau_{x,v}$ corresponds to applying a circular permutation to the angular indexing of k . We define geodesic convolution of a template K and a function f by:

$$(f \circledast k)_v := \max_k \sum_{ij} f_w E_{vijw} K_{i,(j+k) \bmod N_\theta}$$

Likewise directional convolution of an angular function φ by K is defined by:

$$(\varphi \star K)_{vl} = \sum_{ij} \varphi_{wk} \bar{E}_{vijwk} \cdot K_{i,(j+l) \bmod N_\theta}.$$

We also note that the above definitions extend to the case of multi-dimensional filters and matrix-valued templates. Remark that only the triangles supporting the window points will contribute to the result of the convolution, while the value of the signal at the central point will not be taken into account. To achieve this, we add the result of a dense layer to the convolution. That is we multiply the signal at the central point (and direction for dir-conv) by the some fixed (learned) matrix and add the resulting vector to the result of the convolution at every point.

4.5 Spatial Pooling

In traditional CNNs pooling layers refer to a way of sampling the signal. They reduce the resolution of the image by mapping groups of pixels of the original image to a single pixel of a reduced image. The advantage is two-fold: first it allows to summarize the information over a group of pixels and to achieve robustness to local perturbations of the signal and second it reduces computation time and space. A common option called max-pooling is to separate the original image into consecutive square blocks of pixels, where each block is mapped to a single pixel of the reduced image by taking the maximal response over the block in each channel. A closely

related notion, which is easily generalizable to meshes is the notion of strided convolution (Figure 6). It consists of spacing the pixels of the filter window by $D - 1$ pixels and applying it every D pixels thus reducing the output resolution by a factor D .

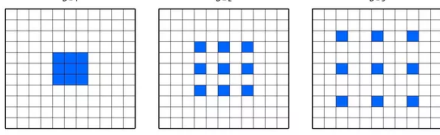


Fig. 6. Strided convolution on images with convolution kernel in blue. Stride of 0 (left), 1 (middle) and 2 (right).

We can see strided convolution as a regular convolution applied to the sub-sampled image. For meshes we define a similar notion by transferring the signal to a coarser mesh and then applying geodesic convolution on the new mesh. In practice, we simplify the original mesh using the classic quadratic edge collapse approach [Garland and Heckbert 1997], which also produces a mapping between the original and the simplified meshes. We use this to transfer both a signal from the original to the simplified mesh (pooling), or from the simplified mesh to the original (un-pooling) by simply picking the value stored at the closest vertex, among those collapsed to it. To define pooling for directional (angular) functions we need to also transfer angles from the original mesh to its simplified counterpart. To transfer local angles from a mesh M to a mesh N given a map from M to N we first compute the 3D rotation sending the normal at each vertex to the normal at its image, this allows us to compare the reference directions on both shapes in the tangent plane to N and to deduce the oriented angular offset between them. We then simply transfer the angles by adding the offset to all discretized angles and interpolating the result between consecutive discrete bins, similarly to our construction of parallel transport.

5 GPC AND PARALLEL TRANSPORT

To compute GPC at each vertex of the mesh we used the algorithm proposed by [Melv er and Reimers 2012] which is a variation of the fast marching algorithm [Sethian 1999] which allows to compute the angular coordinate as well as the geodesic distance. We extended the algorithm of [Melv er and Reimers 2012] to also compute the parallel transport of angles. Fast marching-like algorithms allow to propagate information along meshes, and rely on a local transfer subroutine. A vertex is selected among a set of candidates based on a priority criterion, the information stored at the vertex is then propagated to some of its neighbors based on an update criterion using the local transfer subroutine. The algorithm stops once a certain final condition is met. In our case a vertex is updated until the radius exceeds a certain threshold $R_{\max} > 0$. In practice, we follow the basic approach of [Melv er and Reimers 2012] for propagating information, but in addition to updating the geodesic distance ρ and polar angle θ we also keep track of the difference in angles between the reference directions at the source and target points. The original algorithm [Melv er and Reimers 2012] has a subroutine for updating the GPC angle at a vertex inside a triangle given estimates at the two other vertices. We use the same subroutine to update the

transported angle given estimates at the two other vertices. Since the representation of directions is different at each vertex, we must transfer estimates to every new vertex. We transfer the angles along the edges connecting the vertices. To transport an angle from vertex i to a neighbor j along the edge e_{ij} we first apply rotations to the GPCs of i and j so that e_{ij} gives the reference direction at i and e_{ji} gives reference direction at j . The rotated GPCs at i and j have a relative angular offset of π which allows to deduce the angular offset between the original GPCs. We transfer the angle at i to j along e_{ij} by adding the angular offset between the GPCs at i and j . We use the same edge transfer to transfer the reference direction at the source to its neighbors in order to initialize the algorithm candidates set. This modified version allows us to compute the angle difference between the initial and final reference directions, which in turn provides the estimate for the parallel transport of the unit directions along geodesics between points.

Figure 7 illustrates a single GPC window and parallel transport on a sphere. Namely, Figure 7a illustrates the angular coordinates of the GPC window via color coding, and the parallel transport of a particular direction from the center vertex to other vertices in the window, whereas Figure 7b illustrates our discretization of the angular and radial bins.

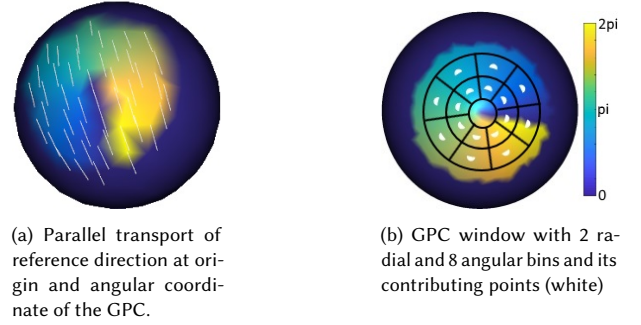


Fig. 7. An example of parallel transport and the GPC window on a sphere. The colors represent the angular coordinate in the GPC window.

6 EVALUATION

6.1 Architecture

We implemented our deep learning pipeline with Keras [Chollet et al. 2015] using Tensorflow backend [Abadi et al. 2015]. We based our architecture on Residual networks [He et al. 2016]. For classification tasks we used an architecture organized in stacks of ResNet blocks (Figure 8). Each stack is the composition of a fixed number of ResNet blocks. After each stack the mesh is down-sampled by a factor 4 using a pooling layer. The radius and number of filters is multiplied by two to preserve time and space complexity across stacks (Figure 9). For classification tasks we apply angular max pooling and average the signal over the shape we then apply softmax classifier on the resulting vector. For segmentation tasks we need to produce a point-wise prediction therefore the signal needs to be up-sampled back to the original shape. We combined our ResNet architecture with U-net [Ronneberger et al. 2015]. Our U-ResNet architecture (Figure 10) consists of two blocks an encode and a decode block. The encode

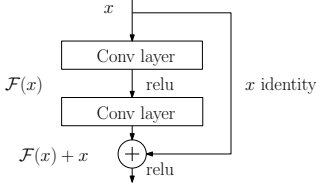


Fig. 8. ResNet block

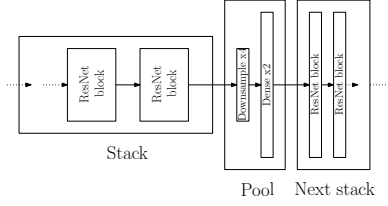


Fig. 9. Our ResNet architecture

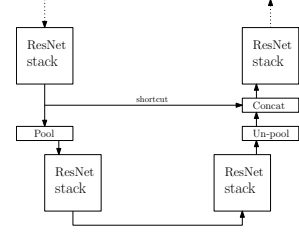


Fig. 10. U-ResNet architecture

block is a copy of our ResNet architecture, the decode block is similar but pooling is replaced by un-pooling. After each stack the window radius is divided by 2, the signal dimension is divided by two and its up-sampled by 4. Shortcut connection are added to help keeping spatial localization information:

6.2 Experiments

In our experiments we compared MDGCNN to GCNN [Masci et al. 2015] and PointCNN [Li et al. 2018] in image classification and to GCNN, PointNet++ [Qi et al. 2017] and Dynamic Graph CNN [Wang et al. 2017] in shape segmentation and shape matching tasks using various input features. We trained all networks using ADAM [Kingma and Ba 2014] optimizer with learning rate 0.001. Since MDGCNN and GCNN are closely related we used the same architectures and window radii for both to make the comparison as fair as possible. For experiments with varying domains we first center the shapes and normalize them so that they have unit variance, then we use a fixed initial radius for the whole dataset. The memory complexity is linear in the number of vertices, the number of radial bins and the number of directional bins of the windows. MDGCNN and GCNN have similar time and memory complexity but MDGCNN uses more complex tensor indexing to align local windows and performs directional interpolation of the result of convolution. Therefore our implementation of MDGCNN is slightly slower in practice. For example in our image classification experiment on 50000 images mapped to spheres with 3000 vertexes we observed times of 7 min 13s for one epoch with GCNN and 10 min 13 with MDGCNN using a GTX 1080 graphics card. This speed advantage however is compensated by a the significantly faster convergence of MDGCNN as shown in Figure 14.

6.3 Image classification

In our first experiment we compare MDGCNN, GCNN and PointCNN on the CIFAR-10 image classification benchmark [Krizhevsky and Hinton 2009] on different domains. We do not try to achieve state-of-the-art performance. Our purpose is to demonstrate that MDGCNN is able to learn complex signals over mesh domains and is superior to GCNN. The CIFAR-10 dataset consists of 60000 32 by 32 RGB images in 10 classes, (airplane, automobile, bird, cat, deer, dog, frog, horse, ship, truck) with 6000 images per class. There are 50000 training images and 10000 test images. We mapped the images to two different meshes, a regular grid and a sphere. In the case of the sphere, we parametrized two opposite hemispheres in polar

Table 1. Classification accuracy on the CIFAR-10 dataset using different methods on different domains.

CIFAR 10 - classification		
Method	Domain	Accuracy
GCNN	sphere	0.6712
MDGCNN	sphere	0.7706
GCNN	grid	0.6767
MDGCNN	grid	0.7932
PointCNN	grid	0.7669

coordinates and then used elliptical mapping [Fong 2015] (shown in Figure 11) to map the images to both hemispheres.

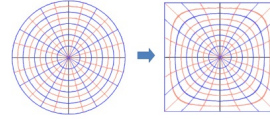


Fig. 11. Elliptical mapping from a disc to a square

Once the mapping is computed every mesh vertex is equipped with 2D coordinates which allows us to pull back the images using bilinear interpolation inside pixels. The resulting image on the sphere is then linearly interpolated inside each triangle as shown on Figure 12. Let us note that in both the grid and the sphere case, using principal curvature directions to fix the reference orientation, as done in [Boscaini et al. 2016; Monti et al. 2017], would not be meaningful as every point is an umbilic on these domains, so that every direction is a principal one.



Fig. 12. Examples of CIFAR-10 images mapped to a sphere using elliptical mapping

We trained our ResNet architecture (Figure 9) in batch of 10 images with 3 ResNet stacks, one ResNet block (Figure 8) per stack and 16 filters for the first stack, we chose an initial radius of 1.8 pixels, that is 1.8 in the grid case and $(32 \times 1.8) / \pi$ on the sphere. The network converged after 50 epochs. The results scores in Table 1 show a clear advantage for MDGCNN over GCNN, we also observe similar scores

for MDGCNN on different domains despite the important stretching introduced by the elliptical mapping of images to the sphere and the irregularity of the sphere meshing compared to a grid mesh. In their recent work [Li et al. 2018] Li et al. applied PointCNN to the CIFAR-10 classification benchmark. The results summarized in Table 1 suggest that our method compares favorably to PointCNN in the image classification context.

6.4 Shape segmentation

In our second experiment we compared our MDGCNN against GCNN and several other state-of-the-art methods: Toric cover CNN [Maron et al. 2017], PointNet++ [Qi et al. 2017] and Dynamic graph CNN [Wang et al. 2018]. We evaluated all methods on the human segmentation benchmark proposed by [Maron et al. 2017]. This dataset consists of 370 models from SCAPE, FAUST, MIT and Adobe Fuse [Adobe 2016]. All models are manually segmented into eight labels, three for the legs, two for the arms, one for the body and one for the head. The test set is the 18 models from the SHREC07 dataset in human category.

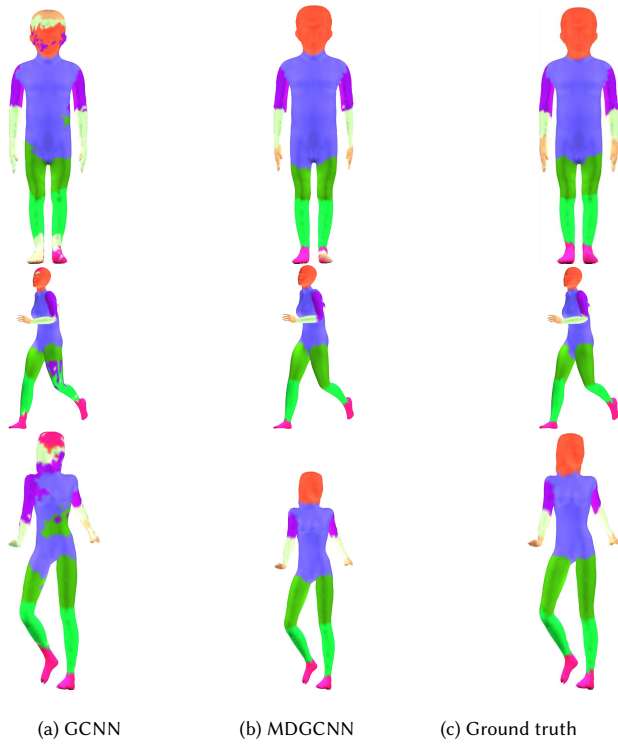


Fig. 13. Human shapes segmentation comparison between standard GCNN and MDGCNN (ours) using the 3D coordinates of the shapes as input. The data is augmented by random rotations and scaling at training to ensure rigid motion invariance and improve robustness of the learning.

Table 2. Segmentation accuracy of several methods on the human body dataset introduced in [Maron et al. 2017].

Human body segmentation				
Method	Input feat.	# feat.	epochs	accuracy
Toric cover	WKS, AGD, curv.	20+2+4	20	0.88
Pointnet++	3D coords	3	200	0.9077
DynGraphCNN	3D coords	3	200	0.8972
GCNN	3D coords	3	200	0.7649
MDGCNN	3D coords	3	50	0.8861
GCNN	WKS, curv.	20+4	50	0.8489
MDGCNN	WKS, curv.	20+4	50	0.8612
GCNN	SHOT ₆	64	50	0.3888
MDGCNN	SHOT ₆	64	50	0.8530
GCNN	SHOT ₉	64	50	0.7410
MDGCNN	SHOT ₉	64	50	0.8879
GCNN	SHOT ₁₂	64	50	0.8640
MDGCNN	SHOT ₁₂	64	50	0.8947

Table 3. Standard deviation of test accuracy of MDGCNN and GCNN across 5 independent runs.

Human body segmentation			
method	Input	epochs	standard deviation
MDGCNN	3D coords	50	0.01059
GCNN	3D coords	50	0.11487
MDGCNN	3D coords	100	0.02831
GCNN	3D coords	100	0.08007
MDGCNN	3D coords	200	0.00454
GCNN	3D coords	200	0.05513

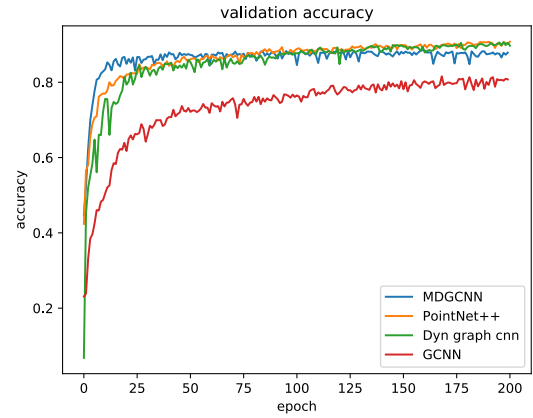


Fig. 14. Validation accuracy per epoch on the human body segmentation benchmark introduced in [Maron et al. 2017] using global 3D coordinates as input.

We used our U-ResNet architecture with 2×2 of two blocks, 16 filters on the first layer and an initial radius of 0.1. Since we did not have access to the same features as used in [Maron et al.

2017], we used different inputs, SHOT [Salti et al. 2014] and WKS [Aubry et al. 2011] descriptors as well as the 3D coordinates of the shape, we denote by SHOT_k the 64-dimensional SHOT descriptor with window radius equal to k percent of the shape area. For the experiments taking the 3D coordinates of the shape as input we first apply a random rotation and scaling between 0.85 and 1.15 to learn features that are robust to global transformations. Table 2 summarizes the scores obtained by different methods.

In addition to improving accuracy, we have also observed that training MDGCNN can be significantly more stable compared to GCNN. In Table 3 we report the standard deviation of the validation accuracy of MDGCNN and GCNN across 5 independent runs with different number of epochs. Here the training data is fixed and the variance is only due to the stochastic nature of the optimization procedure. We observe in Figure 13 that MDGCNN is better than GCNN at learning features that are invariant to rigid motion directly from the 3D coordinates of the shape without *a priori* descriptors assumed in the data, simply via data augmentation. Learning global features from the 3D coordinates of the shape vertices requires aggregating information between possibly distant points. Since directional convolution allows better communication between distant points, MDGCNN noticeably outperforms GCNN when using 3D coordinates as input, producing much smoother results. This is also illustrated in the qualitative results shown in Figure 13. On the other hand, shape descriptors often carry more global information about the points for example SHOT relies on local histograms counting the mesh vertexes and normals into bins, while WKS relies on diffusion processes along the surface. We observe close performance in favor of MDGCNN when using WKS. For SHOT we notice that GCNN performance considerably degrades when reducing the radius of the SHOT windows while MDGCNN is able to maintain its performance much better. PointNet++ and Dynamic Graph CNN behaved similarly on this benchmark we observed slightly slower convergence compared to MDGCNN but slightly better final results (see Figure 14). We observe similar accuracy between MDGCNN and the Toric Cover method of [Maron et al. 2017]. We note, however, that the Toric Cover method is based on non-canonical mappings from the torus to the surface and requires considerable data augmentation by examining many such mappings. This results in long training and prediction computation times. According to [Maron et al. 2017] it takes about 5 hours for the Toric Cover method to complete one epoch at training using 6 Nvidia K80 GPUs. Using 50 different mappings, it takes 45 minutes to calculate predictions on the human class of SHREC07 while it takes 1 min 10s to train our MDGCNN network on one epoch using a single Nvidia TITAN Xp card and 2.2317s to calculate predictions on the test set.

6.5 Shape matching

Finally, we also applied our pipeline in the context of non-rigid shape matching on the FAUST dataset, used in [Masci et al. 2015]. In this experiment, goal is to predict the index corresponding to each vertex in the 0-th shape of the dataset. The original experiment in [Masci et al. 2015] used the GCNN architecture using SHOT descriptors as input. In order to remove the bias present in the data, due to all meshes sharing the same connectivity, we also re-meshed the

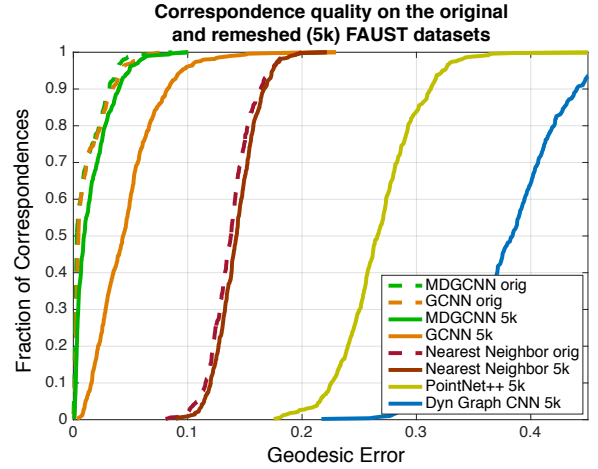


Fig. 15. Performance of shape correspondence on the FAUST dataset and its re-meshed version (5k) evaluated by plotting the fraction of correspondences within a geodesic radius of the ground truth. Higher curve corresponds to better performance.

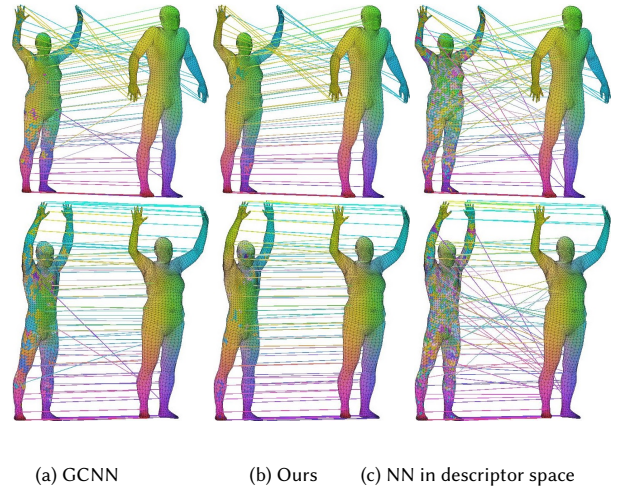


Fig. 16. Shape correspondence on the a remeshed (5k) version of the FAUST dataset using : (a) GCNN [Masci et al. 2015] (b) Our method MDGCNN, and (c) using the nearest neighbor in the SHOT descriptors space.

FAUST shapes from 6890 vertexes to 5000 vertexes to evaluate the robustness of both algorithms. We used our U-ResNet architecture with 2×2 stacks of two blocks with 16 filters on the first layer and an initial radius of 0.1 taking SHOT_{12} as input. The network was trained for 100 epochs for MDGCNN and 200 for GCNN. We measured the geodesic error between the predicted labels and the ground truth on the 0-th shape for both the original FAUST shapes and the re-meshed ones (5k) (Figure 15).

Contrary to [Masci et al. 2015] we do not use post processing on the predictions of the algorithm, and measure the accuracy directly

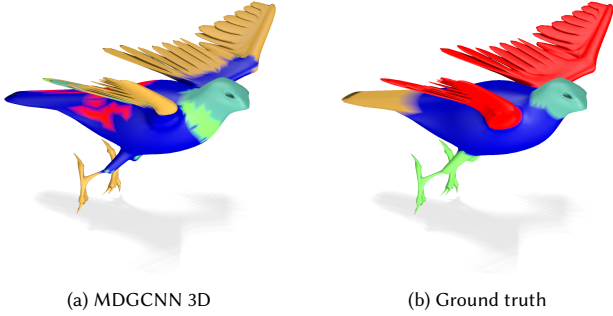


Fig. 17. Limiting case of purely intrinsic pipelines

on the output of the networks. As a baseline we also measured the geodesic error of correspondences obtained on 100 random pairs of the test set using nearest neighbors in the space of SHOT descriptors. We see that learning techniques vastly outperform the baseline. On the original set with shapes having the same connectivity MDGCNN and GCNN behave similarly. On the re-meshed set GCNN noticeably degrades while MDGCNN is able to maintain its precision. Figure 16 shows several examples of correspondences between pairs of shapes computed using different methods.

Figure 15 also shows a comparison with PointNet++ [Qi et al. 2017] and Dynamic graph CNN [Wang et al. 2018] using 3D coordinates as input (trained for 200 epochs). Note that these methods are designed for segmentation tasks and are not suited for shape matching in their current form.

6.6 Limitations & Future Work

Our pipeline still has important limitations. The learning process depends on the construction of local coordinate systems which might not be suited to describe certain types of patterns possibly introducing a bottleneck to the learning. More specifically constructions such as geodesic polar coordinates and parallel transport are purely intrinsic based on the metric of the surface, therefore some areas that are close in the embedding space might be considered far in this representation. A typical limiting case of purely intrinsic pipelines such as ours is when some region of the shape is made of multiple parts that seem to merge in a single object, they might very well fail to recognize it as as a such. We illustrate this by applying our segmentation pipeline to the bird class of PSB dataset, shown in Figure 17. From a purely topological perspective the bird’s wings are locally disconnected as they are made of many feathers.

As we can see in Figure 17 our pipeline most likely recognized each independent feather as a bird tail, failing to recognize the wing in its entirety. Another limiting case of pipelines based on local coordinates is the presence of very thin nearly degenerate parts as the bird’s legs and claws. Two dimensional systems of coordinates might not be appropriate to model near one dimensional parts. On the other hand, Figure 18 shows that in the absence of such limiting cases a bird shape can be properly segmented by MDGCNN.

Another limitation shared by CNNs for image processing is the choice of scale (size) of local windows. The same features can have different meaning depending on the scale at which they are detected.

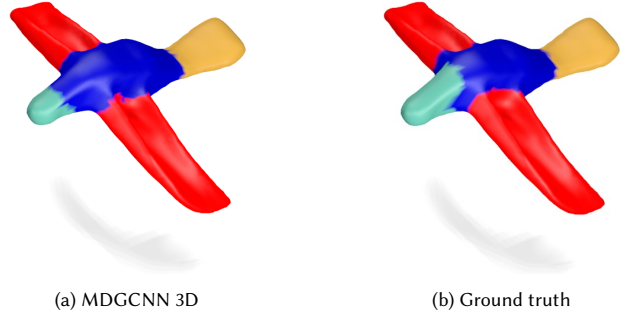


Fig. 18. Successful segmentation with our approach.

Since our current implementation only uses a single fixed scale, this may limit the generalization power of our method in situations where relative proportions of object parts are different from the examples in the training set. In image processing the scale issue has been addressed e.g. by using the notion of Inception modules [Szegedy et al. 2015]. We leave implementation of inception modules within the MDGCNN framework for future work.

Our discretization also has some issues. For example, the parallel transport of a direction might fall between angular bins attached to the target vertex. For this reason we used linear interpolation between the two adjacent directions. However this still introduces a non-negligible error which grows with the number of layers especially for low number of angular bins. A possible alternative is to use the Fourier basis to represent directional signals at all points. This would allow exact direction transfer since rotations act linearly on the basis functions. However activation functions and operations such as angular max pooling would be harder to perform.

Perhaps the most immediate, and relatively straightforward extension of our work would be to use our multi directional approach in the context of local parameterizations via anisotropic diffusion kernels [Boscaini et al. 2016], but without assuming a canonical reference direction at every point. More generally, it would be interesting to use multi-scale approaches with different ways of computing the coordinate systems and transporting the information depending on the scale possibly using extrinsic information to help the network learning different semantic interpretations across different scales in order to improve the overall robustness. Finally other constructions of directional convolution with potentially stronger properties are possible, and we give an example in the Appendix Section 8.2.

7 CONCLUSION

In this work we presented a novel approach to define convolution over curved surfaces that do not admit global or canonical coordinate systems. Namely, we proposed a way to align *local* systems of coordinates allowing to build and learn consistent filters that can then be naturally used across different domains. Our approach is built on the notion of directional functions, which generalize real-valued signals. We proposed a technique to convolve such directional functions with learned template (or filter) functions to produce new directional functions. This allows us to compose these

convolution operations, without any loss of directional information across layers of a neural network.

We showed that our new approach compares favorably to its most direct analogues producing smoother and more robust results and can compete with more recent techniques, even when using “weak” input signals such as the 3D coordinates of the points. We believe that idea of multi-directional convolution can be generalized and can open the door to addressing many other situations where the data representation is ambiguous by allowing a neuron to have different ways of interpreting its input and to communicate with its contributors. For example, in our case a neuron can process its input depending on the choice of reference direction and propagate this information to its contributors. However, other types of alignment between layers can be thought of across different types of ambiguities in the signal as well.

ACKNOWLEDGMENTS

The authors would like to thank the anonymous reviewers for their valuable comments and helpful suggestions. Parts of this work were supported by a Google Focused Research Award and the ERC Starting Grant No. 758800 (EXPROTEA).

REFERENCES

- Martin Abadi, Ashish Agarwal, Paul Barham, and others. 2015. TensorFlow: Large-Scale Machine Learning on Heterogeneous Systems. (2015). <https://www.tensorflow.org/> Software available from tensorflow.org.
- Adobe. 2016. Adobe Fuse 3D Characters. (2016). <https://www.mixamo.com>
- Mathieu Aubry, Ulrich Schlickewei, and Daniel Cremers. 2011. The wave kernel signature: A quantum mechanical approach to shape analysis. In *ICCV Workshops*. IEEE, 1626–1633.
- Davide Boscaini, Jonathan Masci, Simone Melzi, Michael M Bronstein, Umberto Castellani, and Pierre Vandergheynst. 2015. Learning class-specific descriptors for deformable shapes using localized spectral convolutional networks. In *Computer Graphics Forum*, Vol. 34. Wiley Online Library, 13–23.
- Davide Boscaini, Jonathan Masci, Emanuele Rodola, and Michael M. Bronstein. 2016. Learning shape correspondence with anisotropic convolutional neural networks. In *arXiv:1605.06437*.
- Michael M Bronstein, Joan Bruna, Yann LeCun, Arthur Szlam, and Pierre Vandergheynst. 2017. Geometric deep learning: going beyond euclidean data. *IEEE Signal Processing Magazine* 34, 4 (2017), 18–42.
- François Chollet and others. 2015. Keras. <https://github.com/fchollet/keras>. (2015).
- Michaël Defferrard, Xavier Bresson, and Pierre Vandergheynst. 2016. Convolutional neural networks on graphs with fast localized spectral filtering. In *Advances in Neural Information Processing Systems*. 3844–3852.
- Danielle Ezuz, Justin Solomon, Vladimir G Kim, and Mirela Ben-Chen. 2017. GWCNN: A Metric Alignment Layer for Deep Shape Analysis. In *Computer Graphics Forum*, Vol. 36. Wiley Online Library, 49–57.
- Chamberlain Fong. 2015. Analytical methods for squaring the disc. *arXiv preprint arXiv:1509.06344* (2015).
- Michael Garland and Paul S. Heckbert. 1997. Surface Simplification Using Quadric Error Metrics. In *Proceedings of the 24th Annual Conference on Computer Graphics and Interactive Techniques (SIGGRAPH '97)*. 209–216.
- Paul Guerrero, Yanir Kleiman, Maks Ovsjanikov, and Niloy J Mitra. 2018. PCPNet Learning Local Shape Properties from Raw Point Clouds. In *Computer Graphics Forum*, Vol. 37. Wiley Online Library, 75–85.
- Kaiming He, Xiangyu Zhang, Shaoqing Ren, and Jian Sun. 2016. Deep residual learning for image recognition. In *Proceedings of the IEEE conference on computer vision and pattern recognition*. 770–778.
- Evangelos Kalogerakis, Melinos Averkiou, Subhransu Maji, and Siddhartha Chaudhuri. 2017. 3D Shape Segmentation with Projective Convolutional Networks. In *Proc. CVPR*.
- Diederik P Kingma and Jimmy Ba. 2014. Adam: A method for stochastic optimization. *arXiv preprint arXiv:1412.6980* (2014).
- Roman Klokov and Victor Lempitsky. 2017. Escape from cells: Deep kd-networks for the recognition of 3d point cloud models. In *2017 IEEE International Conference on Computer Vision (ICCV)*. IEEE, 863–872.
- I. Kostrikov, Z. Jiang, D. Panozzo, D. Zorin, and J. Bruna. 2017. Surface Networks. *ArXiv e-prints* (May 2017).
- Alex Krizhevsky and Geoffrey Hinton. 2009. Learning multiple layers of features from tiny images. (2009).
- Alex Krizhevsky, Ilya Sutskever, and Geoffrey E Hinton. 2012. Imagenet classification with deep convolutional neural networks. In *Advances in neural information processing systems*. 1097–1105.
- Yangyan Li, Rui Bu, Mingchao Sun, and Baoquan Chen. 2018. PointCNN. *arXiv preprint arXiv:1801.07791* (2018).
- Haggai Maron, Meirav Galun, Noam Aigerman, Miri Trope, Nadav Dym, Ersin Yumer, Vladimir G Kim, and Yaron Lipman. 2017. Convolutional Neural Networks on Surfaces via Seamless Toric Covers. *SIGGRAPH*.
- Jonathan Masci, Davide Boscaini, Michael M. Bronstein, and Pierre Vandergheynst. 2015. Geodesic convolutional neural networks on Riemannian manifolds. In *Proc. of the IEEE International Conference on Computer Vision (ICCV) Workshops*. 37–45.
- Daniel Maturana and Sebastian Scherer. 2015. Voxnet: A 3d convolutional neural network for real-time object recognition. In *Intelligent Robots and Systems (IROS), 2015 IEEE/RSJ International Conference on*. IEEE, 922–928.
- Eivind Lyche Melvæ and Martin Reimers. 2012. Geodesic polar coordinates on polygonal meshes. In *Computer Graphics Forum*, Vol. 31. Wiley Online Library, 2423–2435.
- Federico Monti, Davide Boscaini, Jonathan Masci, Emanuele Rodola, Jan Svoboda, and Michael M. Bronstein. 2017. Geometric Deep Learning on Graphs and Manifolds Using Mixture Model CNNs. In *CVPR*. IEEE Computer Society, 5425–5434.
- Maks Ovsjanikov, Mirela Ben-Chen, Justin Solomon, Adrian Butscher, and Leonidas Guibas. 2012. Functional maps: a flexible representation of maps between shapes. *ACM Transactions on Graphics (TOG)* 31, 4 (2012), 30.
- Charles R Qi, Hao Su, Kaichun Mo, and Leonidas J Guibas. 2017. Pointnet: Deep learning on point sets for 3d classification and segmentation. *Proc. Computer Vision and Pattern Recognition (CVPR), IEEE* 1, 2 (2017), 4.
- Charles R Qi, Hao Su, Matthias Nießner, Angela Dai, Mengyuan Yan, and Leonidas J Guibas. 2016. Volumetric and multi-view cnns for object classification on 3d data. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*. 5648–5656.
- Charles Ruizhongtai Qi, Li Yi, Hao Su, and Leonidas J Guibas. 2017. Pointnet++: Deep hierarchical feature learning on point sets in a metric space. In *Advances in Neural Information Processing Systems*. 5105–5114.
- Olaf Ronneberger, Philipp Fischer, and Thomas Brox. 2015. U-net: Convolutional networks for biomedical image segmentation. In *International Conference on Medical image computing and computer-assisted intervention*. Springer, 234–241.
- Samuele Salti, Federico Tombari, and Luigi Di Stefano. 2014. SHOT: Unique signatures of histograms for surface and texture description. *Computer Vision and Image Understanding* 125 (2014), 251–264.
- James Albert Sethian. 1999. *Level set methods and fast marching methods: evolving interfaces in computational geometry, fluid mechanics, computer vision, and materials science*. Vol. 3. Cambridge university press.
- Konstantinos Sfikas, Theoharis Theoharis, and Ioannis Pratikakis. 2017. Exploiting the PANORAMA representation for convolutional neural network classification and retrieval. In *Eurographics Workshop on 3D Object Retrieval*.
- Baoguang Shi, Song Bai, Zhichao Zhou, and Xiang Bai. 2015. Deeppano: Deep panoramic representation for 3-d shape recognition. *IEEE Signal Processing Letters* 22, 12 (2015), 2339–2343.
- Ayan Sinha, Jing Bai, and Karthik Ramani. 2016. Deep learning 3d shape surfaces using geometry images. In *European Conference on Computer Vision*. Springer, 223–240.
- Justin Solomon, Gabriel Peyré, Vladimir G Kim, and Suvrit Sra. 2016. Entropic metric alignment for correspondence problems. *ACM Transactions on Graphics (TOG)* 35, 4 (2016), 72.
- Hang Su, Subhransu Maji, Evangelos Kalogerakis, and Erik Learned-Miller. 2015. Multi-view convolutional neural networks for 3d shape recognition. In *Proceedings of the IEEE international conference on computer vision*. 945–953.
- Christian Szegedy, Wei Liu, Yangqing Jia, Pierre Sermanet, Scott Reed, Dragomir Anguelov, Dumitru Erhan, Vincent Vanhoucke, and Andrew Rabinovich. 2015. Going deeper with convolutions. In *Proc. CVPR*. 1–9.
- Peng-Shuai Wang, Yang Liu, Yu-Xiao Guo, Chun-Yu Sun, and Xin Tong. 2017. O-cnn: Octree-based convolutional neural networks for 3d shape analysis. *ACM Transactions on Graphics (TOG)* 36, 4 (2017), 72.
- Yue Wang, Yongbin Sun, Ziwei Liu, Sanjay E. Sarma, Michael M Bronstein, and Justin M Solomon. 2018. Dynamic graph CNN for learning on point clouds. *arXiv preprint arXiv:1801.07829* (2018).
- Lingyu Wei, Qixing Huang, Duygu Ceylan, Etienne Vouga, and Hao Li. 2016. Dense human body correspondences using convolutional networks. In *Computer Vision and Pattern Recognition (CVPR), 2016 IEEE Conference on*. IEEE, 1544–1553.
- Zhirong Wu, Shuran Song, Aditya Khosla, Fisher Yu, Linguang Zhang, Xiaoou Tang, and Jianxiong Xiao. 2015. 3d shapenets: A deep representation for volumetric shapes. In *Proceedings of the IEEE conference on computer vision and pattern recognition*. 1912–1920.
- Kai Xu, Vladimir G Kim, Qixing Huang, Niloy Mitra, and Evangelos Kalogerakis. 2016. Data-driven shape analysis and processing. In *SIGGRAPH ASIA 2016 Courses*. ACM, 4.

8 APPENDIX

Proof of Proposition 3.1. We have:

$$\begin{aligned} (\overline{\exp}_x^X)^* \tilde{f}(p) &= \tilde{f}(\exp_x^X(p), \Gamma_{x,p}(p/||p||)) \\ &= f(\exp_x^X(p)) := ((\exp_x^X)^* f)(p) \end{aligned}$$

thus:

$$\begin{aligned} \max_{v \in T_x X} (\tilde{f} \star k)(x, v) &= \max_{v \in T_x X} \langle (\overline{\exp}_x^X)^* \tilde{f}, \tau_{x,v}^* k \rangle_{L^2} \\ &= \max_{v \in T_x X} \langle (\exp_x^X)^* f, \tau_{x,v}^* k \rangle_{L^2} \\ &=: f \circledast k(x) \end{aligned}$$

Proof of Proposition 3.2. The first equality holds directly by definition. Namely, by applying the definition of directional convolution of the angular function φ_{Re} with reference direction Re we have:

$$\varphi_{Re} \star k := (\varphi \star k)_{Re}.$$

To prove the second equality observe that Re_x is simply the rotation of e_x by angle θ_x therefore:

$$\begin{aligned} (\varphi \star k)_{Re}(x, \theta) &:= (\varphi \star k)(x, Re_x(\theta)) \\ &= (\varphi \star k)(x, e_x(\theta + \theta_x)) \\ &=: (\varphi \star k)_e(x, \theta + \theta_x) \\ &=: (\varphi_e \star k)(x, \theta + \theta_x) \end{aligned}$$

which proves the proposition using the coordinate free definition of the directional convolution operator.

Below, we provide an alternative proof that shows that Proposition 3.2 still holds when the directional convolution operator $\star$ is defined in the angular coordinate setting, which thus provides a more direct link to the practical setting. The family of unit vectors $e_x \in T_x X$ defines polar coordinate systems on each tangent plane $T_x X$. A tangent vector $p \in T_x X$ is then represented by a tuple (r, θ) where r is its radius and θ is the angle between e_x and $p = re_x(\theta)$ where $e_x(\theta)$ denote the direct rotation of e_x by angle θ . We represent unit vectors only by their angle. Adapting the notations of Eq. (3) in the polar coordinate systems defined by e we denote:

$$\begin{aligned} \tau_{x,\theta}^e &:= \tau_{x,e_x(\theta)} \\ \exp_{x,e_x}^X(r, \theta) &:= \exp_x^X(r, e_x(\theta)) \end{aligned}$$

and $\Gamma_{x,(r,\theta_1)}^e(\theta_2)$ is the angle between $e_{\exp_{x,e_x}^X(r,\theta)}(r,\theta)$ and $\Gamma_{x,Re_x(\theta_1)}(e_x(\theta_2))$ i.e.

$$e_{\exp_{x,e_x}^X(r,\theta)}(\Gamma_{x,(r,\theta_1)}^e(\theta_2)) := \Gamma_{x,Re_x(\theta_1)}(e_x(\theta_2))$$

We first observe that:

$$\begin{aligned} \varphi_{R,e}(x, \theta) &= \varphi_e(x, \theta + \theta_x), \\ \tau_{x,\theta}^{R,e} &= \tau_{x,\theta+\theta_x}^e, \\ \exp_{x,R,e}^X(r, \theta) &= \exp_{x,e_x}^X(r, \theta + \theta_x), \\ \Gamma_{x,\theta_1}^{R,e}(\theta_2) &= \Gamma_{x,\theta_1+\theta_x}^e(\theta_2 + \theta_x) \end{aligned}$$

We have:

$$\begin{aligned} (R,e)_{\exp_{x,R,e}^X(r,\theta_1)}(\Gamma_{x,(r,\theta_1)}^{R,e}(\theta_2)) &= \\ e_{\exp_{x,e_x}^X(r,\theta_1+\theta_x)}(\Gamma_{x,(r,\theta_1+\theta_x)}^e(\theta_2 + \theta_x)) &= \\ R,e_{\exp_{x,R,e}^X(r,\theta_1)}(\Gamma_{x,(r,\theta_1+\theta_x)}^e(\theta_2 + \theta_x) - \theta_{\exp_{x,e_x}^X(r,\theta+\theta_x)}) \end{aligned}$$

Thus:

$$\Gamma_{x,(r,\theta_1)}^{R,e}(\theta_2) = \Gamma_{x,(r,\theta_1+\theta_x)}^e(\theta_2 + \theta_x) - \theta_{\exp_{x,e_x}^X(r,\theta+\theta_x)}$$

Therefore

$$\begin{aligned} (\overline{\exp}_x^{X,R,e})^* \varphi_{R,e}(x, (r, \theta)) &= \\ := \varphi_{R,e}(\exp_{x,R,e}^X(r, \theta), \Gamma_{x,(r,\theta)}^{R,e}(\theta)) &= \\ = \varphi_{R,e}(\exp_{x,e_x}^X(r, \theta + \theta_x), \Gamma_{x,(r,\theta+\theta_x)}^e(\theta + \theta_x) - \theta_{\exp_{x,e_x}^X(r,\theta+\theta_x)}) &= \\ = \varphi_e(\exp_{x,e_x}^X(r, \theta + \theta_x), \Gamma_{x,(r,\theta+\theta_x)}^e(\theta + \theta_x)) &= \\ = (\overline{\exp}_x^{X,e})^* \varphi_e(x, (r, \theta + \theta_x)). \end{aligned}$$

So that:

$$\begin{aligned} (\varphi_{R,e} \star k)(x, \theta) &= \\ = \langle (\overline{\exp}_x^{X,R,e})^* \varphi_{R,e}, (\tau_{x,\theta}^{R,e})^* k \rangle &= \\ = \langle (\overline{\exp}_x^{X,e})^* \varphi_e(\bullet, \bullet + \theta_x), (\tau_{x,\theta+\theta_x}^e)^* k \rangle &= \\ = (\varphi_e \star k)(x, \theta + \theta_x) \end{aligned}$$

8.1 Details on the implementation of convolution

Our practical implementation of geodesic convolution relies on dense tensors for efficiency optimization. We describe it in the following subsection.

Let X be a triangle mesh with N_v vertexes. The exponential map over X is given by geodesic polar coordinates (GPC) around each vertex. We model it as two N_v by N_v matrices r and θ where r_{ij} and θ_{ij} represent the radius and angle at vertex j of the GPC centered at vertex i . The associated Euclidean coordinates $(x_{ij}, y_{ij})_{ij}$, extend inside triangles by linear interpolation. We use the GPC to construct windows at each vertex along which we can transfer signals on the mesh. The windows are defined by their radius R . The window attached to vertex i consists of the points of polar coordinates are $(\frac{j \cdot R}{N_\rho}, \frac{2k \cdot \pi}{N_\theta})_{j \in [1, N_\rho], k \in [1, N_\theta]}$ in the GPC at i . Each window vertex lies inside a triangle we denote by E_{ijkl} the index of the l -th vertex of the triangle containing p_{ijk} the jk -th point of i -th window and by W_{ijkl} the associated barycentric coordinate. That is:

$$p_{ijk} = \sum_{l=1}^3 W_{ijkl}(x_{i,E_{ijkl}}, y_{i,E_{ijkl}}).$$

A a -dimensional signal f on the mesh consist of a N_v by a matrix, it can be pulled back to the window system by the following formula:

$$E^* f_{ijkl} := \sum_{m=1}^3 W_{ijkm} f_{E_{ijkm}, l}$$

this can be seen as a discretization of the pull back by exponential map. Template functions are stacked in ab -polar kernel tensors of shape (N_ρ, N_θ, a, b) to be convolved with a -dimensional signals. In our context we define the geodesic convolution of a a dimensional

signal f by the a -polar kernel K as the b -dimensional directional signal:

$$(f \circledast K)_{ijk} := \sum_{r,m,l} E^* f_{irm} K_{r,(m+j) \bmod N_{\theta},l,k}$$

We adapt the original definition of geodesic convolutional layer:

DEFINITION 8.1 (GEODESIC CONVOLUTIONAL LAYER). *The geodesic convolutional layer of ab -kernel K , central kernel C and bias vector B and activation function ξ transforms any a -dimensional signal f into the b -dimensional signal:*

$$g_{K,C,B,\xi}^c(f)_{ij} := \max_k \xi(f \circledast K_{ikj} + \sum_l C_{jl} f_{il} + B_j)$$

We model the parallel transport as a N_v by N_v matrix γ where γ_{ij} is the angular offset in the angular coordinates θ of the GPC induced by the parallel transport along the geodesic joining vertex i to vertex j . We discretize the transport of angles by the 4D tensor Γ of shape $(N_v, N_\rho, N_\theta, 3)$ defined by:

$$\Gamma_{ijkl} := \frac{\gamma_{i,E_{ijkl}}}{2\pi} + \frac{k}{N_\theta}$$

is the normalized angle representing the unit radial vector at jk -vertex of i -th window in its GPC. The tensor Γ stores exact angular values however since discretized directional functions are defined using N_θ evenly spaced angles we consider the lower and upper transport tensors $[\Gamma]$ and $[\Gamma]$. We can interpolate directional signals inside the resulting angular sectors by using the fractional part $\{\Gamma\} := \Gamma - [\Gamma]$. We define the discrete pull back of a directional signal φ and by the "completed" exponential map as the tensor:

$$\varphi(E, \Gamma)_{ijkl} := \sum_m W_{ijkm} ((1 - \{\Gamma\}_{ijkm}) \varphi_{E_{ijkm}, [\Gamma]_{ijkm}, l} + \{\Gamma\}_{ijkm} \varphi_{E_{ijkm}, [\Gamma]_{ijkm}, l})$$

We define the discrete directional geodesic convolution of a a -dimensional directional signal φ by the ab -polar kernel K as the b -dimensional directional signal:

$$\varphi \star K_{ijk} := \sum_{r,m,l} \varphi(E, \Gamma)_{irm} K_{r,(m+j) \bmod N_{\theta},l,k}$$

DEFINITION 8.2 (DIRECTIONAL GEODESIC CONVOLUTIONAL LAYER). *The directional geodesic convolution of ab -polar kernel K , central kernel C and bias vector $B \in \mathbb{R}^b$ and activation function ξ transforms any a -dimensional directional signal φ to the b -dimensional directional signal:*

$$\text{dir}_{K,C,B,\xi}(\varphi)_{ijk} := \xi(f \star K_{ijk} + \sum_l C_{kl} f_{il} + B_k)$$

8.2 A Stronger Notion of Directional Convolution

Several definition of directional convolution are possible we chose to transport tangent directions along geodesic to produce a local radial vector field on the surface which is rotation invariant this simplifies implementation. There is an other natural choice to pull back directional signals to the tangent plane we can set:

$$\overline{\exp}_{x,v}^X(p) = (\exp_x^X(p), \Gamma_{x,p}(v)),$$

and define convolution of a directional function φ over X by a template k by:

$$(\varphi \star_X k)(x, v) = \langle \overline{\exp}_{x,v}^X \rangle^* \varphi, \tau_{x,v}^* k \rangle$$

On the plane ($X = \mathbb{R}^2$) the above formula can be simplified, we have:

$$\overline{\exp}_{x,v}^{\mathbb{R}^2}(p) = (x + p, v)$$

To ease notation we identify $\mathbb{R}^2$ with $\mathbb{C}$. The map τ simply rotates the template:

$$\tau_{x,e^{i\theta}}^* k := k(e^{-i\theta} \bullet)$$

therefore for any directional function φ over $\mathbb{R}^2$:

$$(\varphi \star_{\mathbb{R}^2} k)(x, e^{i\theta}) = (\varphi(\bullet, e^{i\theta}) * k(e^{-i\theta} \bullet))(x)$$

Let f be a function over $\mathbb{R}^2$ denote by $(x, v) \mapsto \tilde{f}(x, v) := f(x)$ the associated directional function then:

$$(\tilde{f} \star_{\mathbb{R}^2} k)(x, e^{i\theta}) = (f * k)(e^{-i\theta} \bullet)(x)$$

The resulting directional function essentially stores the result of the usual convolution with the rotated kernel in each direction. Thanks to the above observation we can show that an image based CNN using this notion of directional convolution is equivalent to taking the maximal response under rotations of the kernels with its standard CNN counterpart as shown by the following proposition. For simplicity we consider one dimensional signals although the property still holds in higher dimensions.

PROPOSITION 8.1. *Let f a function over $\mathbb{R}^2$, $(K_l)_l$ a sequence of template kernels, $(b_l)_l$ a sequence of real numbers $(\xi_l)_l$ a sequence of activation functions. We define the sequence $(\varphi_l)_l$ of directional functions over $\mathbb{R}^2$ and the sequence $(f_l^\theta)_l$ of function over $\mathbb{R}^2$ by:*

$$\begin{cases} \varphi_0 = \tilde{f} \\ \varphi_{l+1} = \xi_l \circ (\varphi_l \star_{\mathbb{R}^2} K_l + b_l) \end{cases} \quad \begin{cases} f_0^\theta = f \\ f_{l+1}^\theta = \xi_l \circ ((f_l^\theta * K_l)(e^{-i\theta} \bullet) + b_l) \end{cases}$$

Then for all n we have:

$$\varphi_n(x, e^{i\theta}) = f_n^\theta(x).$$

PROOF. We proceed by recurrence. The property is true for $n = 0$ by definition, suppose it is true for $n = k$. We have:

$$\begin{aligned} \varphi_{k+1}(x, e^{i\theta}) &= \xi_k(\varphi_k \star_{\mathbb{R}^2} K_k(x, e^{i\theta}) + b_k) \\ &= \xi_k((\varphi_k(\bullet, e^{i\theta}) * K_k)(e^{-i\theta} \bullet)(x) + b_k) \\ &= \xi_k((f_k^\theta * K_k)(e^{-i\theta} \bullet)(x) + b_k) \\ &:= f_{k+1}^\theta(x) \end{aligned}$$

which proves the property for $n = k + 1$ and concludes the proof. $\square$