

Farthest Sampling Segmentation of Triangulated Surfaces

V. Hernández Mederos 
Instituto de Cibernética,
Matemática y Física

D. Martínez 
Universidad Federal
do Amazonas

J. Estrada Sarlabous 
Instituto de Cibernética,
Matemática y Física

V. Guerra Ones 
Universidad de La Laguna

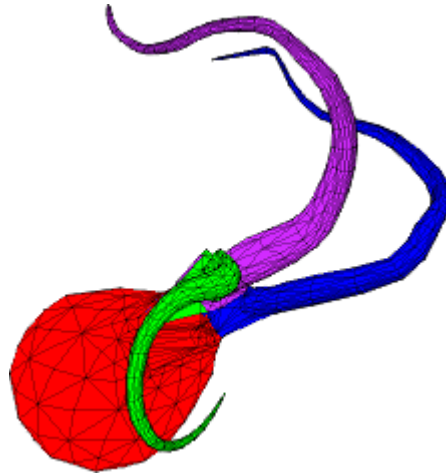


Figure 1. Segmentation of the octopus based on a sample of less than 2% of its triangles.

Abstract

In this paper we present *Farthest Sampling Segmentation* (FSS), a new method for the segmentation of a triangulated surface with n faces into patches. The method is based on the selection of a sample composed by k faces of the triangulation, with $k \ll n$. These faces are chosen by farthest point sampling with respect to a given metric. Pairwise face distances among the faces of the triangulation and the faces in the sample are used to compute an $n \times k$ affinity matrix W^k . Rows of W^k encode the similarity among all triangles and the sample. The segmentation is obtained by applying the clustering algorithm *k-means++* to the rows of W^k . Several theoretical results that support the success of FSS are presented. We explain the

relation between FSS and spectral segmentation methods. Moreover, it is shown that FSS is coherent and stable. An extensive numerical experimentation is included, with several metrics and a large variety of 3D triangular meshes. The quality of the segmentations is measured in terms of Rand and Jaccard distances between FSS and ground-truth segmentations. The results show that always connected clusters are produced and that segmentations obtained when k is less than 10% of n are as good as those obtained using the full affinity matrix. FSS has several advantages. It does not depend on parameters to be tuned by hand and is very flexible, since it can handle any metric. Moreover, it is a very cheap method, with a computational cost of $\mathcal{O}(knm)$, where m is the cost to evaluate a metric between two faces of the triangulation.

1. Introduction

Mesh segmentation is an important ingredient of many geometric processing and computer graphics tasks, such as shape matching, parametrization, mesh editing and compression, texture mapping, morphing, multiresolution modeling, animation, and 3D printing. It explains why this subject has received a lot of attention in recent years. In a review of mesh segmentation techniques, Shamir [2008] formulated the segmentation problem as an optimization problem and considered two qualitatively different types of segmentation: the part type, aiming to partition the surface into volumetric parts, and the surface type, attempting to segment the surface into patches. Segmentation techniques are also classified in correspondence with general clustering algorithms, such as region growing, hierarchical clustering, iterative clustering, spectral analysis, etc.

The most important task concerning shape segmentation is how to define a part of the surface. This is done by using various mesh properties or features such as area, size or length, curvature, geodesic distances, normal directions, distance to the medial axis, and shape diameter. In many segmentation algorithms [de Goes et al. 2008; Katz and Tal 2003; Koschan 2003; Lee et al. 2005; Li and Peng 2020; Liu and Zhang 2004; Wang et al. 2016; Zhang and Liu 2005] part analysis is carried out by using surface-based computations. For instance, Zhang and Liu [2005] combined geodesic and angular distances to define a metric that encodes distances between mesh faces, while de Goes et al. [2008] used diffusion distance to propose a hierarchical segmentation method for articulated bodies. The previous approaches are based on intrinsic metrics on the surfaces and do not capture explicit volumetric information. Liu et al. [2009] defined a volumetric part-aware metric combining the volume enclosed by the surface with geodesic and angular distances. This metric is successfully applied in various applications including mesh segmentation, shape registration, part-aware sampling, and shape retrieval.

According to the underlying technique, many segmentation algorithms [de Goes

et al. 2008; Katz et al. 2005; Liu et al. 2006; Liu and Zhang 2004, 2007; Zhang and Liu 2005] belong to the class of the so-called *spectral methods*. In these algorithms, an affinity or Laplacian matrix is constructed by using intrinsic metrics. The original surface is projected into low-dimensional spaces, which are derived from the eigenvectors of the affinity or Laplacian matrix. As a consequence of the Polarization Theorem [Brand and Huang 2003], higher-quality cut boundaries may be obtained from these embeddings. For details about the spectral approach for mesh processing and analysis, including mesh compression, correspondence, parameterization, segmentation, surface reconstruction, and remeshing, see the excellent survey by Zhang et al. [2010].

1.1. Contributions

The main contribution of this paper is a new algorithm for segmentation of triangulated surfaces based on the selection of a sample composed by k distinguished faces. Given a metric to measure the distance between neighboring faces, the matrix W^k encoding the affinity among all faces and the faces in the sample is computed. The segmentation is performed by applying a classical clustering algorithm to the rows of W^k . The new method, called *Farthest Sampling Segmentation* (FSS), does not require to compute the spectrum of the affinity matrix W or of any of its submatrices. Hence, it is computationally cheaper than the segmentation algorithms based on eigendecompositions. See Figure 2 for a graphical overview of the method.

From the theoretical point of view, our first result is the proof that given a sample W^k of the columns of W , the orthogonal projection of W in the space generated by the columns of W^k is the same as the orthogonal projection of W in the space generated by the largest eigenvectors of W , approximated using Nyström’s method for the same sample. This theoretical result clarifies the point of contact between our method and the spectral approach and explains the success of FSS. Moreover, it is shown that if the columns of W^k correspond to the k farthest triangles in the selected

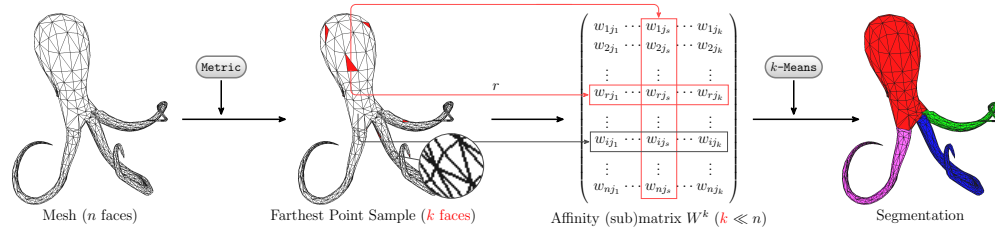


Figure 2. Overview of the FSS algorithm. A small sample of k mesh faces (depicted in red) is selected by means of farthest point sampling with respect to a suitable metric. Columns of the $n \times k$ affinity matrix W^k corresponding to the sample are computed. Mesh segmentation is obtained clustering the k -dimensional rows of W^k .

metric, then for increasing size k , pairs of faces that are close in the selected metric project in pairs of rows of (order $n \times k$) matrix W^k that are close as points in $\mathbb{R}^k$, and also pairs of faces that are far away in the selected metric project in pairs of rows of W^k that are far away as points in $\mathbb{R}^k$ (Lemma 1). This holds independently of the choice of the first element of the sample (Lemma 2).

A wide experimentation illustrating the quantitative and qualitative performance of the FSS method is also included. These experiments make apparent the robustness and approximation power of the algorithm for small samples. Furthermore, a deterministic criterion is proposed to choose a lower bound for the size k of the sample that furnishes a good approximation to W .

1.2. Paper Organization and Notation

In Section 2 we introduce the basic concepts that allow to define the similarity between any two triangles of the mesh. Two low-dimensional embeddings of the affinity matrix are described in Section 3: the spectral and the statistical leverage. The first theoretical result of the paper is included in that section. The embedding based on the computation of the columns of the affinity matrix that corresponds to the farthest triangles is introduced in Section 4, where the coherence and stability of the FSS method is shown and the associated algorithm is explained. Section 5 starts with an experiment to prove that the embedding based on the farthest triangles provides a good approximation of W . Moreover, several numerical experiments are included showing the quantitative and qualitative performance of FSS. A comparison of FSS with the spectral method is finally presented. The last section concludes the paper.

We use capital letters to denote matrices and the same lowercase letter to denote its elements. For example the element i, j of the matrix A is denoted by a_{ij} . Moreover, $A_{i\cdot}$ and $A_{\cdot j}$ represents the i -th row and the j -th column of matrix A respectively, while A^+ denotes the Moore–Penrose inverse of matrix A . All mesh segmentations shown in this work are computed without including any procedure to improve the smoothness of the boundaries of the segments or their concavity, such as proposed by Shapira et al. [2008] and Wang et al. [2014].

2. Distance and Affinity Matrices

Denoting by T the triangulation composed by a set $F = \{f_1, f_2, \dots, f_n\}$ of faces, the segmentation problem consists in defining a partition of F . In most of the segmentation algorithms, an important step to group the elements of F consists in introducing pairwise face distances and constructing an affinity matrix by using them. In the literature, several face distances have been considered [Gal et al. 2007; Katz and Tal 2003; Liu et al. 2009; Shapira et al. 2008]. For instance, Katz and Tal [2003] defined the distance between two adjacent faces as a convex combination of their geodesic and

angular distances. Other metrics have been specially designed to capture parts of the volume enclosed by the surface, such as the part-aware distance [Liu et al. 2009] and the shape-diameter function (SDF) [Shapira et al. 2008]. The part-aware metric by Liu et al. [2009] happens to be expensive, since its computation requires performing two samplings of the triangular mesh by using ray-shooting. On the other hand, as remarked by Liu et al. [2009], the SDF does not capture well the volumetric context.

2.1. Distance Matrix

Denote by f_i and f_j two triangles of T sharing an edge. Assume that we have already defined a distance d_{ij} between faces f_i and f_j . For instance d_{ij} could be the angular distance [Katz and Tal 2003] defined as $\eta(1 - \langle n_i, n_j \rangle)$, where $\langle n_i, n_j \rangle$ denotes the scalar product between the normalized normal vectors n_i and n_j to the triangles f_i and f_j , respectively, and η is a weight introduced to reinforce the concavity of the angles. Another distance very common in the literature is the geodesic distance [Katz and Tal 2003], which in the case of neighboring triangles is defined as the length of the shortest path between their barycenters b_i and b_j . Our third test distance, introduced by Liu et al. [2009], is based on a scalar function defined on the triangulation, the shape-diameter function [Shapira et al. 2008]. The SDF distance between any two adjacent faces f_i and f_j of T is defined as $|SDF(b_i) - SDF(b_j)|$. Each of these test distances captures different features of the triangulation.

The distance d_{ij} between any pair of faces f_i and f_j is computed using the weighted dual graph G_d of T . The i -th node of G_d represents the triangle f_i in T for $i = 1, \dots, n$, and there is an edge between the i -th and the j -th nodes of G_d if the faces f_i and f_j share an edge on the triangulation T . The weight of the edge joining neighboring faces f_i and f_j is defined as $\max\{d_{ij}, \varepsilon\}$, where $\varepsilon > 0$ is very small. The distance is extended to non-neighboring faces as the length of the shortest path between their corresponding nodes in G_d , which may be computed using Dijkstra's algorithm. Observe that the distance d_{ij} satisfies the axioms of metric. We denote by $D = (d_{ij}), i, j = 1, \dots, n$, the matrix of the distances between each pair of faces of the triangulation.

2.2. Affinity Matrix

Given a suitable metric that allows to compute the pairwise distance between faces of the triangulation, the affinity matrix W encodes the probability of each pair of faces being part of the same cluster and can be considered as the adjacency matrix of the weighted graph G_d previously introduced.

Assume that the distance d_{ij} between any pair of faces f_i and f_j of the triangulated surface has been already computed. Then the affinity w_{ij} between faces f_i and f_j , which are closer, should be large. In the literature it is customary to use a Gaussian

kernel to define w_{ij} , $i, j = 1, \dots, n$, as

$$w_{ij} = e^{-d_{ij}/(2\sigma^2)}, \quad (1)$$

where $\sigma = \frac{1}{n^2} \sum_i \sum_j d_{ij}$. Observe that $0 < w_{ij} \leq 1$ and $w_{ii} = 1$ for all $i = 1, \dots, n$. Moreover, $W = (w_{ij})$, $i, j = 1, \dots, n$, is a symmetric matrix. Denote by M the diagonal matrix

$$M = \text{diag}(m_{ii}), \quad (2)$$

where $m_{ii} = \sum_{j=1}^n w_{ij}$.

Many papers in the literature deal with normalized versions of the affinity matrix, which are called Laplacians in the more general context of clustering the data for exploratory analysis [von Luxburg 2007]. For instance, the (nonsymmetric) affinity matrix $M^{-1}W$ is used in spectral image segmentation [Shi and Malik 1997], while the symmetric normalized affinity matrix $Q = M^{-1/2}WM^{-1/2}$ is used for data clustering [Liu and Zhang 2004] and mesh segmentation [Ng et al. 2001]. In applications, the affinity matrix W of order n is huge, therefore segmentation methods requiring the computation of all matrix entries are very expensive. To overcome this problem, the segmentation algorithm proposed in this paper computes only a few columns of matrix W .

3. Low-Dimensional Embeddings for Clustering

In the geometric processing community, low-dimensional embeddings are frequently used to transform the input data from its original domain to another domain. The main purpose of these embeddings is to reduce the dimensionality of the problem, preserving the information of the original data in such a way that the solution of the new problem is cheaper and easier. The segmentation problem can also be considered as a clustering problem, which is frequently solved by applying the *k-means* method [Lloyd 1982]. In this context, dimensionality reduction for *k-means* is strongly connected with a low-rank approximation of the matrix containing the data to be clustered [Boutsidis et al. 2015]. In our problem, the matrix containing the information about “data points” is the affinity matrix W . Each row of W represents the affinity between a triangular face and the rest of the faces. Hence, a valid strategy to solve the segmentation problem consists in computing a low-rank approximation of the affinity matrix W and clustering its rows.

3.1. Spectral Approach

The most popular low-dimensional embeddings in the literature are the spectral ones, which are constructed from a set of eigenvectors of a properly defined linear operator. They have been successfully applied in mesh segmentation [de Goes et al. 2008; Liu et al. 2006; Liu and Zhang 2004, 2007] and also in other geometric processing

applications, such as shape correspondence [Jain et al. 2007] and retrieval [Elad and Kimmel 2003] and mesh parametrization [Gotsman 2003; Mullen et al. 2008].

From the theoretical point of view, spectral embeddings are supported by a classical linear algebra result, the Eckart–Young theorem [Eckart and Young 1936]. It establishes that the best rank- k approximation in the Frobenius norm of a real, symmetric, and positive semi-definite matrix W of dimension n is the matrix

$$E^k = \tilde{U}^k (\tilde{U}^k)^\top, \quad (3)$$

where $\tilde{U}^k$ is the matrix with columns $\sqrt{\lambda_1}u_1, \sqrt{\lambda_2}u_2, \dots, \sqrt{\lambda_k}u_k$ and $u_1, u_2, \dots, u_k$ are the eigenvectors of W corresponding to its largest eigenvalues $\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_k$. It means that the following equality holds for the Frobenius norm of the error:

$$\|W - E^k\|_F = \min_{\substack{X \in \mathbb{R}^{n \times k} \\ \text{rank}(X) \leq k}} \|W - XX^\top\|_F. \quad (4)$$

If we denote by U^k the matrix with columns $u_1, u_2, \dots, u_k$, then $\tilde{U}^k = (\Lambda^k)^{\frac{1}{2}} U^k$, where Λ^k is the order- k diagonal matrix with diagonal elements $\lambda_1, \lambda_2, \dots, \lambda_k$. Moreover, it is not difficult to prove that $(U^k)^+ = (U^k)^\top$. Hence, the orthogonal projection $U^k(U^k)^+W$ of W on the space generated by columns of U^k satisfies

$$\begin{aligned} U^k(U^k)^+W &= U^k(U^k)^\top W &&= U^k(\Lambda^k)(U^k)^\top \\ &= \left(\tilde{U}^k(\Lambda^k)^{-\frac{1}{2}} \right) (\Lambda^k) \left(\tilde{U}^k(\Lambda^k)^{-\frac{1}{2}} \right)^\top \\ &= \tilde{U}^k(\tilde{U}^k)^\top &&= E^k. \end{aligned}$$

The equality $E^k = U^k(U^k)^+W$ means that the best rank- k approximation of W is the projection of W on the space generated by the eigenvectors of W corresponding to its largest eigenvalues.

Spectral clustering algorithms also rely on the Polarization Theorem [Brand and Huang 2003], which states that as the dimensionality of the spectral embeddings decreases, the clusters in the data are better defined. In practical applications, it is necessary to choose a value of k representing a good compromise between these two apparently conflicting results. This value should be small enough to obtain a good polarization of the embedding data, but at the same time large enough to reduce the distortion of the proximity among the data due to the embedding.

Spectral methods of segmentation are in general expensive, since they require the computation of eigenvalues and eigenvectors of the Laplacian matrix. In some cases [Liu et al. 2006; Liu and Zhang 2004], the Laplacian matrix is obtained, introducing a normalization of the affinity matrix. In other cases, it arises from a discretization of the Laplace–Beltrami operator [Reuter et al. 2009]. In the geometry processing context the Laplacian used in segmentation is a dense and usually very

large matrix. To face this problem, Liu et al. [2006] used Nyström's method, since it only requires a small number of sampled rows of the affinity matrix and the solution of a small-scale eigenvalue problem.

More precisely, a sample $\mathcal{X}$ of k mesh faces define a partition $F = \mathcal{X} \cup \mathcal{Y}$ of F . Let $p_{\mathcal{X}}$ be the set of k indices of the faces in the sample $\mathcal{X}$ and $p_{\mathcal{Y}}$ the set of $n - k$ indices of the faces in $\mathcal{Y}$. Denote P as the permutation matrix corresponding to the vector $p = (p_{\mathcal{X}}, p_{\mathcal{Y}})$. Then the permuted affinity matrix $PWP^{\top}$ has the following structure:

$$PWP^{\top} = \begin{bmatrix} A & B \\ B^{\top} & C \end{bmatrix}, \quad (5)$$

where A is the order- k affinity matrix of the elements in $\mathcal{X}$ and B is the order- $(k \times (n - k))$ matrix of the cross-affinities between elements in $\mathcal{X}$ and $\mathcal{Y}$. The eigenvectors of W corresponding to the k largest eigenvalues, i.e., the columns of U^k , may be approximated [Fowlkes et al. 2004; Liu et al. 2006] by the columns of the $n \times k$ matrix $P^{\top} N^k$, with

$$N^k = \begin{bmatrix} U_A \\ B^{\top} U_A \Lambda_A^{-1} \end{bmatrix}, \quad (6)$$

where $A = U_A \Lambda_A U_A^{\top}$ is the spectral decomposition of A . The orthogonal projection F^k of W on the space generated by the columns of matrix $P^{\top} N^k$ provided by Nyström's method is given by

$$F^k = (P^{\top} N^k) (P^{\top} N^k)^+ W. \quad (7)$$

The accuracy of the eigenvectors computed by using Nyström's method strongly depends on the selection of the k columns of W corresponding to sample $\mathcal{X}$. Therefore, different schemes have been considered in the literature, for instance random sampling, uniform sampling, max-min farthest point sampling, and greedy sampling [Fowlkes et al. 2004; Kumar et al. 2012; Liu et al. 2006; Mahoney 2011; de Silva and Tenenbaum 2004].

Proposition 1. *Let W be a symmetric order- n matrix and $p = (p_{\mathcal{X}}, p_{\mathcal{Y}})$ a permutation vector of indices $1, 2, \dots, n$, where $p_{\mathcal{X}}$ has size k . Let P be the order- n permutation matrix represented by vector p and W^k the $n \times k$ matrix whose columns are the columns of W with indices in $p_{\mathcal{X}}$. Then, if W^k has full rank, it holds that*

$$(W^k) (W^k)^+ = (P^{\top} N^k) (P^{\top} N^k)^+, \quad (8)$$

where N^k is given by Equation (6).

Proof. Since columns of W^k are the columns of W with indices in $p_{\mathcal{X}}$, from Equation (5) we get,

$$PW^k = \begin{bmatrix} A \\ B^{\top} \end{bmatrix}. \quad (9)$$

By the hypothesis W^k has full rank, hence from Equation (9) it holds that A and Λ_A are nonsingular. Moreover, from $A = U_A \Lambda_A U_A^\top$ it follows that $U_A = A U_A \Lambda_A^{-1}$. Hence, using Equation (6) and Equation (9), we get

$$N^k = \begin{bmatrix} A U_A \Lambda_A^{-1} \\ B^\top U_A \Lambda_A^{-1} \end{bmatrix} = \begin{bmatrix} A \\ B^\top \end{bmatrix} U_A \Lambda_A^{-1} = P W^k U_A \Lambda_A^{-1}. \quad (10)$$

But columns of W^k and $W^k U_A$ are linearly independent, thus

$$(W^k U_A \Lambda_A^{-1})^+ = \Lambda_A U_A^+ (W^k)^+. \quad (11)$$

Finally, from Equations (10) and (11), it holds that

$$(P^\top N^k)(P^\top N^k)^+ = (W^k U_A \Lambda_A^{-1})(W^k U_A \Lambda_A^{-1})^+ = (W^k)(W^k)^+.$$

□

Remarks. From the previous result it holds that:

1. The orthogonal projection H^k of W on the space generated by the columns of W^k given by

$$H^k = W^k (W^k)^+ W \quad (12)$$

coincides with the orthogonal projection F^k of W on the space generated by the columns of matrix $P^\top N^k$, which is provided by Nyström's method.

2. Since $H^k = F^k$, the approximation to W provided by H^k and the approximation F^k obtained by Nyström's approach have the same accuracy. This accuracy depends on the selection of the k columns of W corresponding to sample $\mathcal{X}$.
3. If we associate the i -th face of T with the i -th row of W , then clustering the rows of W may be replaced by clustering either the rows of W^k or the rows of $P^\top N^k$.

3.2. Statistical Leverage Approach

Mahoney [2011] proposed a method to select a sample of the columns of a matrix W of dimension n in such a way that the space generated by the selected columns provides a good approximation of W . Given k , with $k \ll n$, the method assigns to the j -th column of W a *leverage* or *importance score* π_j that measures the influence of that column in the best rank- k approximation of W .

The use of the leverage scores for column subset selection dates back to 1972 [Jolliffe 1972]. However, the introduction of the randomized approach [Holodnak et al. 2015] has given essential theoretical support to the leverage scores in the role of revealing important information hidden in the underlying matrix structure.

More precisely, if $v_1, v_2, \dots, v_k$ are the right singular vectors of W corresponding to the largest singular values, the leverage π_j is defined as

$$\pi_j = \frac{1}{k} \sum_{i=1}^k (v_i^j)^2, \quad (13)$$

where v_i^j denotes the j -th component of v_i . The normalization factor $\frac{1}{k}$ is introduced in Equation (13) to consider π_j as a probability associated to the j -th column of W . By using that score as an importance sampling probability distribution, the algorithm constructs an $n \times m$ matrix C , composed by $m \geq k$ columns of W [Mahoney 2011]. With high probability, the selected columns are those that exert a large influence on the best rank- k approximation of W . In our experiments in Section 5, we use a slight modification of Mahoney’s algorithm to obtain a matrix C , which we denote by C^k , that is composed exactly by k columns. More precisely, if we arrange in decreasing order the leverages $\pi_{j_1} \geq \pi_{j_2} \geq \dots \geq \pi_{j_n}$, then the i -th column of matrix C^k is the column j_i of W , for $i = 1, \dots, k$. The orthogonal projection of W on the space generated by the columns of C^k is given by

$$G^k = C^k (C^k)^+ W. \quad (14)$$

4. Farthest Sampling Mesh Segmentation Method (FSS)

Computing distances from every node to a subset of nodes of a graph (landmarks or reference objects) is a well-known method to efficiently provide estimates of the actual distance. In this context, this distance information is also referred to as an *embedding*. Landmarks have been used for graph measurements in many applications, such as round-trip propagation, transmission delay, or social search in networks. The selection of the optimal set of landmarks is an NP-hard problem [Kamoussi et al. 2016; Potamias et al. 2009], hence heuristic solutions should be used. The extensive experimentation with several heuristics for the landmark selection on real-world massive graphs presented by Potamias et al. [2009] indicates that a smart landmark selection strategy provides good approximations of the distances in comparison with random landmark selection [Kleinberg et al. 2004]. The analysis of the stretch factor of distance approximations, obtained with landmarks computed with the farthest point sampling strategy, provides some evidence supporting the success of the farthest point sampling as a landmark selection strategy in isometry-invariant shape processing [Kamoussi et al. 2016].

The previous ideas and the result and discussion at the end of Section 3.1 inspired us to propose a mesh segmentation method based on the computation of a small sample of columns of the affinity matrix W . It is well known that the quality of the approximation to W obtained from the eigenvectors computed with Nyström’s method

strongly depends on the selection of the sample. Consequently, our segmentation method FSS consists of the following steps: First we use the farthest point heuristic to select a sample of the faces. Then we compute a matrix W^k encoding the affinity between all faces and the faces in the sample. Finally, a clustering method is applied to the rows of W^k .

The rationality behind these steps is the following. A representative subset of the columns of W must have maximal rank. Since the entries of the j -th column of W are the affinity values between all faces of T and the j -th face, the sample should not include columns corresponding to faces that are very close between them. In this sense, the farthest point heuristic is a good strategy to avoid redundancy in the sample. On the other hand, if two faces f_i and f_j of T are close in the selected metric d , i.e., d_{ij} is small, then the corresponding rows $D_{i,\cdot}$ and $D_{j,\cdot}$ of the full distance matrix D are approximately equal, since according to the triangular inequality the difference between the r -th components of $D_{i,\cdot}$ and of $D_{j,\cdot}$ is bounded by d_{ij} for $r = 1, \dots, n$. Hence, due the continuity of the Gaussian kernel, the rows $W_{i,\cdot}$ and $W_{j,\cdot}$ of the affinity matrix W are also approximately equal. In other words, the proximity among rows of W , and consequently between the faces of T , is well reflected by the proximity of the same rows of W^k .

4.1. Sampling Procedure

Our algorithm FSS is deterministic and greedy in the sense that at each iterative step, it makes a decision about which column to add according to a rule that depends on the already-selected columns. As mentioned before, the sample of W is derived from a sample of columns of the distance matrix D . To obtain a good approximation of D , it is enough to select a set $\mathcal{X} \subset F$ of distinguished faces that can be considered as landmarks, in the sense that the distance between any pair of faces f_i and f_j can be approximated in terms of the distance of f_i (respectively f_j) to the landmark faces.

The method iteratively computes the columns of a matrix X , which contains a sample of k columns of the distance matrix D . More precisely, in the first step we choose randomly a value j_1 with $1 \leq j_1 \leq n$ and define the first column of matrix X as the vector built up with the distances of all faces to the j_1 -th face, $x_{i1} = d_{ij_1}$, $i = 1, \dots, n$. Then, we search the index j_2 of the face farthest from the j_1 -th face and assign to the second column of X the vector of the distances of all faces to the face j_2 . In general, in the step l , for $k \geq l \geq 2$, we have the $l - 1$ indices $j_1, j_2, \dots, j_{l-1}$ previously selected and a matrix X of order $n \times (l - 1)$ with columns $D_{\cdot,j_1}, D_{\cdot,j_2}, \dots, D_{\cdot,j_{l-1}}$, which contains the distances of all faces f_i , $i = 1, \dots, n$ to the faces $f_{j_1}, f_{j_2}, \dots, f_{j_{l-1}}$. In this step we look for the index j_l of the face that maximizes the minimal distance to the faces $f_{j_1}, f_{j_2}, \dots, f_{j_{l-1}}$:

$$j_l = \arg \max_{1 \leq i \leq n} \left\{ \min_{1 \leq r \leq l-1} x_{ir} \right\}, \quad (15)$$

where $x_{ir} = d_{ij_r}$, $i = 1, \dots, n$, $r = 1, \dots, l - 1$ is the element of X in the position (i, r) , i.e., the distance between faces f_i and f_{j_r} . Once we have j_l , we compute the l -th column of X as the vector of distances of all faces to the j_l -th face.

The i -th row of X contains the coordinates of a point in $\mathbb{R}^k$, which could be considered as a k -dimensional embedding of the point in $\mathbb{R}^n$ given by the i -th row of D (which represents the distances of all faces to the face f_i). Given k , by using the matrix X we compute the matrix $W^k = (w_{ij}^k)$ of order $n \times k$, which is an approximation of the submatrix of W composed by the k columns $j_1, \dots, j_k$ in the sample:

$$w_{ij}^k = e^{-x_{ij}/(2\sigma_k^2)}, \quad i = 1, \dots, n, \quad j = 1, \dots, k, \quad (16)$$

where

$$\sigma_k = \frac{1}{n k} \sum_{i=1}^n \sum_{j=1}^k x_{ij}. \quad (17)$$

Finally, it remains to explain how we compute the size k of the sample. In this sense, several options are possible. The simplest one is to define a priori the value of k , for instance as the integer part of a prescribed percent of the total number of faces n . In this case, the *Sampling procedure* is summarized in the procedure in Listing 1.

Another option for computing the size k of the sample is the following. A value $\beta_l \geq 0$, $l \geq 1$, strongly related to the selection of index j_l in Equation (15) is introduced, defining

$$\beta_l := \max_{1 \leq i \leq n} \left\{ \min_{1 \leq r \leq l} x_{i,r} \right\}. \quad (18)$$

Recall that for all $l \geq 1$, $\beta_l \geq 0$. Furthermore, from Equation (18) it is clear that the sequence $\{\beta_l, 1 \leq l \leq n\}$ is monotonic decreasing with $\beta_n = 0$. In fact, while new faces are included in the sample, the distance of the new face that is simultaneously farthest away from all actual members of the sample decreases. If the sample contains

```

Procedure 1
/* Input */
triangulation T, sample size k

- Choose randomly an index  $j_1$  with  $1 \leq j_1 \leq n$ 
for  $l=1$  to  $k$ 
{
- Using Dijkstra, compute  $X_{:,l} = D_{:,j_l}$ 
- Compute  $j_{l+1} = \arg \max_{1 \leq i \leq n} \{\min_{1 \leq r \leq l} x_{i,r}\}$ 
}
/* Output: sampling distance matrix X */
return X
    
```

Listing 1. Sampling procedure, version 1.

```

Procedure 2
/* Input */
triangulation  $T$ ,  $\varepsilon > 0$ 

- Choose randomly an index  $j_1$  with  $1 \leq j_1 \leq n$ 
- Assign  $l \leftarrow 1$ 
- Assign  $\delta \leftarrow 2\varepsilon$ 
while  $\{\delta > \varepsilon \text{ and } l < n\}$ 
{
    - Using Dijkstra, compute  $X_{:,l} = D_{:,j_l}$ 
    - Compute  $\beta_l = \max_{1 \leq i \leq n} \{\min_{1 \leq r \leq l} x_{i,r}\}$  and  $j_{l+1} = \arg \max_{1 \leq i \leq n} \{\min_{1 \leq r \leq l} x_{i,r}\}$ .
    - Assign  $\delta \leftarrow \frac{\beta_l}{\beta_1}$ 
    - Assign  $l \leftarrow l + 1$ .
}
/* Output: sampling distance matrix  $X^*$  */
return  $X$ 
    
```

Listing 2. Sampling procedure, version 2.

all faces, i.e., if $l = n$, then matrix X is a permutation of the columns of the distance matrix D , and for all $i = 1, \dots, n$ it holds that $\min_{1 \leq r \leq n} x_{i,r} = 0$; thus we get $\beta_n = 0$. Hence, the value of β_k may also be interpreted as a measure of the error introduced when the “original” data in a n -dimensional space are substituted by their k -dimensional embedding. Given an upper bound $\varepsilon > 0$, the size k of the sample may be computed as

$$k = \min_{2 \leq l \leq n} \left\{ l \text{ such that } \frac{\beta_l}{\beta_1} < \varepsilon \right\}. \quad (19)$$

For $1 > \varepsilon > 0$, the value of k computed by using Equation (19) depends on ε , and it is usually much smaller than n . Pseudocode of the previous Sampling procedure is included in Listing 2.

4.2. Coherence and Stability of FSS

Now we are ready to explain why FSS is coherent in the sense that clustering the rows of W^k happens to be consistent with clustering the rows of the full matrix W . That is very important, since it is well known that, in general, points far away may project in very close points, giving rise to non-connected clusters. This is not the case in the FSS embedding and, consequently, no artifacts appear in the segmentation process.

From now on we denote the *farthest point (FP) sample* of size k as the set of indices $\mathcal{J}_X^k := \{j_1, j_2, \dots, j_k\}$ such that j_1 is randomly chosen and, for $2 \leq l \leq k$, j_l is the index of the face maximizing the minimal distance to faces $j_1, j_2, \dots, j_{l-1}$. To $\mathcal{J}_X^k$ is associated the $n \times k$ matrix X , submatrix of D . W^k is the $n \times k$ affinity matrix corresponding to X , which is computed according to Equations (16) and (17).

Lemma 1. *Given a triangulation T and a selected metric d , for fixed initial face index j_1 , let $\mathcal{J}_\chi^k$ be the FP sample of size k . Then, there exists k^* such that, if two faces of T are far away (very close, respectively) in the metric d , then for all $k > k^*$, the corresponding rows of matrix W^k are also far away (very close, respectively) as points in $\mathbb{R}^k$.*

Proof. First we prove that for any sample, it holds that pairs of faces that are close in the selected metric d project in pairs of rows of the associated order- $(n \times k)$ matrix W^k that are close as points in $\mathbb{R}^k$. Indeed, if two faces f_i and f_j of T are close in the selected metric d , i.e., if d_{ij} is small, then the corresponding row vectors $D_{i,\cdot}$ and $D_{j,\cdot}$ of the full distance matrix D are approximately equal, since according to the triangular inequality the difference between the r -th components of $D_{i,\cdot}$ and $D_{j,\cdot}$ is bounded above by d_{ij} , i.e.,

$$|d_{ir} - d_{jr}| \leq d_{ij} \quad \text{for } r = 1, \dots, n.$$

Obviously, the same upper bound holds for any row vectors $X_{i,\cdot}$ and $X_{j,\cdot}$ sampled from the $n \times k$ matrix X . Hence, pairs of faces that are close in the selected metric d project in pairs of rows of the order- $(n \times k)$ matrix W^k that are close as points in $\mathbb{R}^k$, independently of the choice of sample.

Alternatively, assume that faces f_i and f_j of T are far away in the selected metric d , i.e., d_{ij} is large. Then for any ε with $d_{ij} \gg \varepsilon > 0$, there exists k^* such that $\beta_k < \frac{\varepsilon}{2}$ for all $k > k^*$. Denote $i^* = \arg \min_{r \in \mathcal{J}_\chi^k} d_{ir}$ and $j^* = \arg \min_{r \in \mathcal{J}_\chi^k} d_{jr}$. From the definition of β_k , it is clear that $d_{ii^*} \leq \beta_k$ and $d_{jj^*} \leq \beta_k$.

Assume that $d_{ii^*} \leq d_{jj^*}$, then $d_{ii^*} \leq d_{ji^*}$. Moreover, by triangular inequality,

$$d_{ji^*} + d_{ii^*} \geq d_{ji}.$$

From the previous inequalities we obtain that, for $r \in \mathcal{J}_\chi^k$, the maximum difference between the r -th components of $D_{i,\cdot}$ and $D_{j,\cdot}$ is bounded below by

$$\max_{r \in \mathcal{J}_\chi^k} |d_{jr} - d_{ir}| \geq |d_{ji^*} - d_{ii^*}| = d_{ji^*} - d_{ii^*} \geq d_{ji} - 2d_{ii^*} \geq d_{ji} - 2\beta_k > d_{ji} - \varepsilon.$$

On the other hand, if $d_{ii^*} \geq d_{jj^*}$, then proceeding in a similar way it holds that

$$\max_{r \in \mathcal{J}_\chi^k} |d_{jr} - d_{ir}| \geq |d_{jj^*} - d_{ij^*}| = d_{ij^*} - d_{jj^*} \geq d_{ji} - 2d_{jj^*} \geq d_{ji} - 2\beta_k > d_{ji} - \varepsilon.$$

Thus, for $k > k^*$, rows i and j of matrix W^k are far away, i.e., pairs of faces that are far away in the selected metric d project in pairs of rows of matrix W^k associated with the FP sample $\mathcal{J}_\chi^k$ that are far away as points in $\mathbb{R}^k$. $\square$

Our experiments in the next section show that it is possible to choose a value of k^* in Lemma 1 such that k^* is large enough to reduce the distortion of the affinity among the data due to the embedding and small enough to obtain a good polarization.

The first step of the Sampling procedure (see Listings 1 and 2) is to randomly select face j_1 . Now we show that choosing at random different indices of the initial triangle, the differences between the i -th and l -th rows of the corresponding affinity matrices associated with the FP samples tend to be equal for increasing size k . Thus, the segmentation result is stable with respect to the selection of the initial face j_1 .

Lemma 2. *Let T be a triangulation, d a selected metric, and $\varepsilon > 0$. Then, there exists k^* such that, for any pair $\mathcal{J}_\mathcal{X}^k$ and $\bar{\mathcal{J}}_\mathcal{X}^k$ of FP samples of size k , $k \geq k^*$, with different initial faces j_1 and $\bar{j}_1$ and associated affinity matrices W^k and $\bar{W}^k$, respectively, it holds that*

$$\left| \|W_{l,\cdot}^k - W_{i,\cdot}^k\| - \|\bar{W}_{l,\cdot}^k - \bar{W}_{i,\cdot}^k\| \right| < \varepsilon$$

for any pair of indices $l, i \in \{1, 2, \dots, n\}$.

Proof. Denote by β_k and $\bar{\beta}_k$ the β values in Equation (18) corresponding to $\mathcal{J}_\mathcal{X}^k = \{j_1, \dots, j_k\}$ and $\bar{\mathcal{J}}_\mathcal{X}^k = \{\bar{j}_1, \dots, \bar{j}_k\}$, respectively. Set $\tilde{\beta}_k = \max\{\beta_k, \bar{\beta}_k\}$. Given $\varepsilon > 0$, if k is sufficiently large, we may assume that $\tilde{\beta}_k < \frac{\varepsilon}{2}$. Then, for $1 \leq r \leq k$, it holds that $d_{ij_r} \leq \tilde{\beta}_k$ and $d_{l\bar{j}_r} \leq \tilde{\beta}_k$ with $1 \leq l, i \leq n$. Hence,

$$|d_{ij_r} - d_{i\bar{j}_r}| \leq \tilde{\beta}_k \quad \text{and} \quad |d_{lj_r} - d_{l\bar{j}_r}| \leq \tilde{\beta}_k.$$

Thus, from the previous inequalities it follows for $1 \leq l, i \leq n$ and $1 \leq r \leq k$ that

$$\begin{aligned} \left| |d_{ij_r} - d_{lj_r}| - |d_{i\bar{j}_r} - d_{l\bar{j}_r}| \right| &\leq |d_{i,j_r} - d_{i,\bar{j}_r} - d_{l,j_r} + d_{l,\bar{j}_r}| \\ &\leq |d_{i,j_r} - d_{i,\bar{j}_r}| + |d_{l,j_r} - d_{l,\bar{j}_r}| \\ &\leq 2\tilde{\beta}_k < \varepsilon. \end{aligned}$$

□

Remarks.

1. If we choose two different initial faces, we obtain in general two different farthest point samples of size k . These samples give rise to different affinity matrices W^k and $\bar{W}^k$. But for k sufficiently large, the differences between any two rows of W^k and the same rows of $\bar{W}^k$ tend to be equal. Since the FSS method constructs the segmentation clustering the rows of W^k and $\bar{W}^k$, it follows from Lemma 2 that for, k sufficiently large, the segmentations obtained starting with different initial faces are the same.
2. In Section 5.4 we show that, choosing a fixed initial triangle j_1 , most of the segmentation results obtained with FP samples of relatively small size k are very close to the segmentation obtained from the full affinity matrix.

```

Procedure FSS
/* Input */
    triangulation  $T$ , number of clusters  $n_c$ 

- Call to the Sampling procedure to compute the sampling distance
  matrix  $X$ .
- Assign  $k$  as the number of columns of  $X$ .
- Compute the normalization factor  $\sigma_k = \frac{1}{n} \sum_{i=1}^n \sum_{j=1}^k x_{ij}$ 
for  $i=1$  to  $n$ 
{
    for  $j=1$  to  $k$ 
        {- Compute  $w_{ij}^k = e^{-x_{ij}/(2\sigma_k^2)}$  }
    - Compute  $normW_i = \|(w_{i1}^k, \dots, w_{ik}^k)\|_2$ 
    for  $j=1$  to  $k$ 
        {- Compute  $w_{ij}^k = w_{ij}^k / normW_i$  }
}
- Apply  $k$ -means to the  $n$  rows of  $W^k = (w_{ij}^k)$  to obtain  $n_c$  clusters.
/* Construct segmentation vector  $s$  */
for  $i=1$  to  $n$ 
    {- Assign to  $s_i \leftarrow$  the index of the cluster of the face  $f_i$  }
/* Output: segmentation vector  $s$  */
return  $s = (s_1, \dots, s_n)$ 
    
```

Listing 3. FSS algorithm.

4.3. Pseudocode of the FSS Method

As previously pointed out, the FSS method first computes a small sample W^k of columns of W . Based on the identification of the rows of W with the faces of T , the segmentation of the mesh is obtained clustering the rows of W^k . Listing 3 includes pseudocode of the FSS algorithm, which receives as input the triangulation T and the number n_c of desired clusters. The number k of columns of the affinity matrix to be computed depends on the sampling procedure. In any case, it is assumed that $k \ll n$. The FSS algorithm calls to the Sampling procedure, which returns the matrix X . This is the main step of FSS. The l -th column of X is the vector of distances of all faces to the j_l -th face in the selected sample. Observe that the FSS algorithm normalizes the rows of W^k (as suggested by Liu and Zhang [2004]), hence these rows may be interpreted as points on the $(k-1)$ -dimensional sphere $\mathbb{S}^{k-1}$.

For the sake of simplicity, the segmentation is carried out by applying the classic k -means clustering algorithm to these points. Recall that the FSS method is not tied to using any specific clustering algorithm. Since the result of k -means depends on the initial seeding, in our implementation we have used several replications with random starting points for the seeds. The final segmentation is defined by the best solution obtained with this strategy.

Compared with other algorithms reported in the literature, the new algorithm has several advantages. First, it only computes an $n \times k$ matrix W^k . On the other hand, unlike previous works [Liu et al. 2006; Liu and Zhang 2004], the FSS algorithm does not require computing the spectrum of W or the spectrum of any submatrix of W .

The construction of matrix W^k requires the computation of

$$(n-1) + (n-2) + \dots + (n-k) = \frac{k}{2}(2n-k-1) = \mathcal{O}(kn)$$

distances between faces. In comparison, obtaining the full affinity matrix W is much more expensive and would require the computation of

$$(n-1) + (n-2) + \dots + 2 + 1 = 1/2 n(n-1) = \mathcal{O}(n^2)$$

distances between faces.

The total computational cost C_t of the FSS algorithm is the sum of the cost C_k of computing the approximation W^k of the affinity matrix W plus the cost C_s of clustering the rows of W^k in n_c clusters. If m is the cost of computing the distance d between two faces of a given triangulation with n faces (this cost strongly depends of the underlying metric: for instance, if the metric is the geodesic or the angular metric, then $m = \mathcal{O}(n \log n)$), then $C_k = \mathcal{O}(knm)$. Since k , n_c , and the number n_i of Lloyd iterations [Lloyd 1982] are bounded, it holds that $C_s = \mathcal{O}(n_i n_c kn)$. Thus, if only $k \ll n$ columns of the affinity matrix W are computed, then the total cost C_t is dominated by C_k , i.e., $C_t = \mathcal{O}(knm)$. In the numerical experiments of the next section, we show that the embedding based on the farthest triangle provides a good approximation of W and therefore it is useful to perform the segmentation.

5. Numerical Experiments

To prove the performance of the mesh segmentation algorithm proposed in this paper, we wrote three main codes. The first code is the basis for the experiment developed in Section 5.1, which illustrates the advantages and limitations of the low-dimensional embeddings previously considered. This code is also used to compute the β curve in Section 5.5. The second code is an implementation of the FSS algorithm whose results are reported in Sections 5.2, 5.3, and 5.4. The last code is an implementation of the segmentation method proposed by Liu et al. [2006], which applies Nyström's method to approximate the spectral embeddings of faces of the triangulation. This code is the computational basis of the comparison developed in Section 5.6 between the performance of FSS and the spectral segmentation method by Liu et al. [2006]. We recall that the mesh segmentations shown in this section are computed without including any procedure to improve the quality (smoothness of the boundaries of the segments or their concavity), as done in other works [Shapira et al. 2008; Wang et al. 2014].

5.1. Low-Dimensional Embeddings

As we previously mentioned, a valid strategy to solve the segmentation problem is based on computing a low-rank approximation of the affinity matrix W . In the following experiment we compare the approximation power of the low-dimensional embeddings described in Sections 3 and 4. Given a mesh with n triangles, we compute the affinity matrix W of order n given by Equation (1). Furthermore, the projection A^k of W on different spaces for $k = 1, \dots, n$ is also computed. More precisely, given a value of k , four projections are computed, each one obtained when A^k is the matrix E^k given by Equation (3), the matrix F^k given by Equation (7), the matrix G^k given by Equation (14), and the matrix H^k given by Equation (12).

For each approximation A^k , the absolute error in Frobenius norm,

$$\text{error}_{\text{abs}} = \|W - A^k\|_F, \quad (20)$$

is computed and compared to the error in Equation (4) of the best approximation E^k .

Figures 3 and 4 show the results obtained for two 3D triangulations representing an octopus (see Figure 10) and a hand (see Figure 11), models 125.off and 200.off, respectively, of the Princeton Segmentation Benchmark [Chen et al. 2009]. In these examples, the distance matrix D was computed using the geodesic metric to measure the distance between triangles. In Figure 3 we plot the absolute error in Equation (20) as a function of k . The red curve corresponds to the error in Equation (20) obtained when A^k is the matrix E^k of the best rank- k approximation of W . Similarly, the black and blue curves are computed with $A^k = G^k, H^k$, respectively. Since $H^k = F^k$ (see Proposition 1), the curve corresponding to Nyström's projection agrees with the blue curve corresponding to the embedding proposed in this paper. Observe that this curve is the closest to the curve obtained for the embedding corresponding to the best approximation of the affinity matrix.

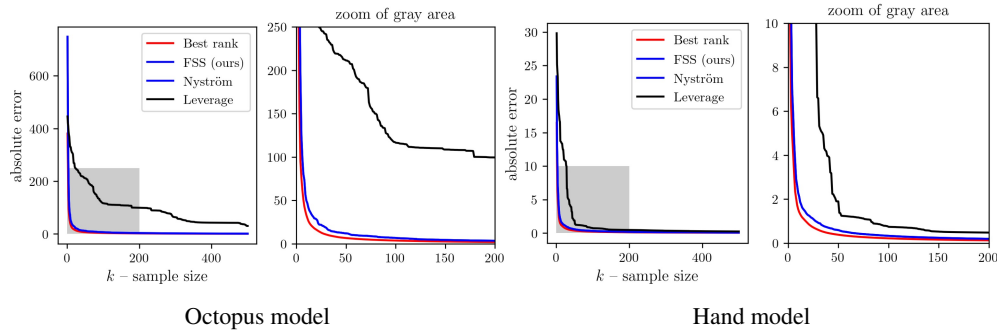


Figure 3. Absolute errors in Equation (20) computed for increasing values of k using different approximations A^k of the affinity matrix W . Left: Octopus model with 2682 faces. Right: Hand model with 3026 faces.

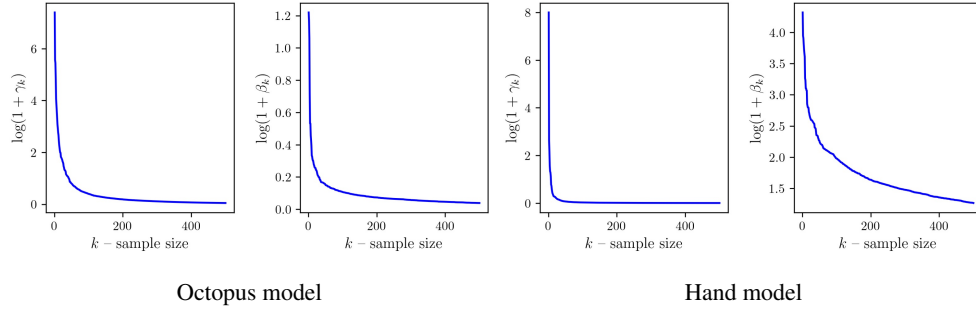


Figure 4. The curves $\log(1 + \gamma_k)$ and $\log(1 + \beta_k)$, both plotted as functions of k for $k = 1, \dots, 500$. Left: The curves for the octopus model with 2682 faces. Right: The curves for the hand model with 3026 faces.

Figure 4 shows the curves $\log(1 + \gamma_k)$ and $\log(1 + \beta_k)$, where γ_k is the k -th singular value of W (in descending order) and β_k is given by Equation (18) (with $l = k$). These curves correspond to the octopus and the hand models and all decrease very fast.

In general, in all our experiments with several triangulations we observed the following:

1. The absolute error curves show that the matrices H^k and F^k (recall that $H^k = F^k$) provide the approximation to W closest to the optimal E^k for $k \leq n/2$. In practice, we are interested in a good rank- k approximation of W with $k \ll n$. Hence, the embedding corresponding to the farthest point sampling provides the better approximation with the lower computational cost. It explains experimentally why the mesh segmentation algorithm FSS compares favorably to other methods reported in the literature, which are based on the spectrum of W and happen to be more expensive.
2. The method proposed in this paper cannot be considered as a random method, since except the first column of the sample, the rest of the columns are selected deterministically. However, one may think of $1 - \beta_l/\beta_1$, $l > 1$ as the conditional probability of selecting the j_l -th column of W given that the columns $j_1, \dots, j_{l-1}$ have been previously selected. In this context, our farthest point sampling scheme may be considered as an algorithm to select the k columns with highest conditional probability (see Section 5.4 for more details).

5.2. Mesh Segmentation Qualitative Performance

In this section we show the performance of our mesh segmentation algorithm FSS with several 3D triangulation models. In the experiments reported here and in the next section, the k -means++ algorithm [Arthur and Vassilvitskii 2007] is applied

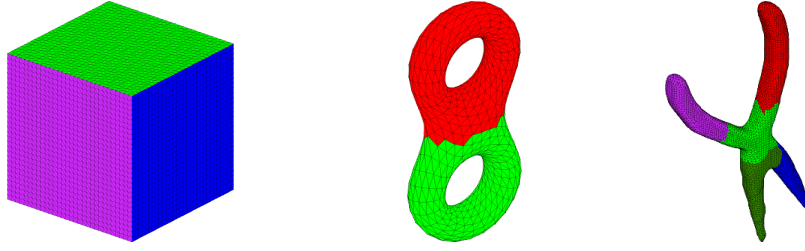


Figure 5. Left: Segmentation with the angular distance of the cube model. The clusters are obtained computing 1% of the 10,800 columns of W . Center: Segmentation of the eight model with the geodesic distance obtained computing 2% of the 1536 columns of W . Right: Pliers model segmentation with the SDF distance obtained computing 5% of the 8970 columns of W .

to the normalized rows of the rectangular matrix W^k , composed by the k selected columns of W . These rows are considered as points in $\mathbb{S}^{k-1}$. To measure the distance between two vectors, we use the cosine distance. Several replicates of the k -means++ algorithm are applied, and for each replicate the seeds of the n_c clusters are selected randomly. The final segmentation is defined by the best solution obtained with this strategy, which typically results in a global minimum of the sum of point-to-centroid distances. In general, the FSS algorithm works very fast since in all segmentations k is at most 10% of the total number of faces. Our goal here is just to evaluate visually the quality of the segmentations produced by the algorithm.

The first example that we considered is the model of a cube defined by a triangulation with 10,800 faces. This model is ideal to check how the algorithm works when the distance between triangles is measured in terms of the angular distance. To segment the model, we computed only 1% of the columns of the affinity matrix W . Figure 5 (left) shows that the results are excellent, since the faces of the cube correspond exactly with the six clusters produced by the automatic segmentation. In the second example, we use the geodesic distance between triangles to define the affinity matrix W of the eight model. In Figure 5 (center) we show the segmentation in two clusters of the model, obtained computing only 2% of the columns of W . Observe that each cluster agrees approximately with one handle of the eight. In our third example, we use the SDF distance to segment the pliers model into five clusters. In Figure 5 (right) we show the results obtained computing 5% of the columns of W . As in the previous examples, the clusters are natural partitions of the model.

In our next examples we use a product metric to compute the distance between neighboring triangles. More precisely, if the faces f_i and f_j share an edge of the triangulation, then the product distance d_{ij} between them is defined as $d_{ij} := d_{ij}^g d_{ij}^a$, where d_{ij}^g and d_{ij}^a are the geodesic and angular distances between f_i and f_j , respectively. As usual, the distance between nonadjacent faces is defined as the length of the shortest path in the dual graph. Figure 6 (left) shows the segmentation in eight

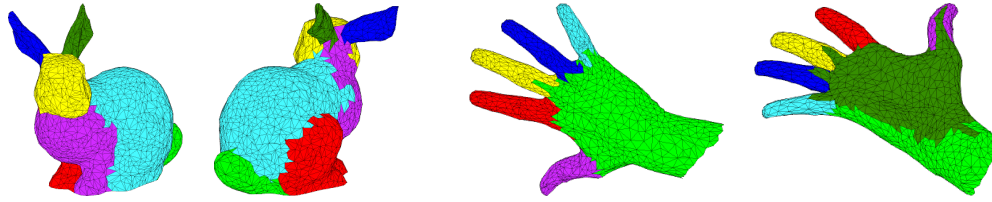


Figure 6. Front and back views of segmentations obtained using the product of geodesic and angular distances as the metric. Left: Bunny model segmented using 10% of the 3860 columns of W . Right: Hand model segmented using 1% of the 3000 columns of W .

clusters of the bunny model. This result was obtained computing 10% of the total number of columns of the affinity matrix W . Observe that the segmentation produced by the product distance distinguishes well not only the big ears but also the small tail. The hand model with 3000 faces is more challenging. As we observe in Figure 6 (right), there is some leakage in the clusters corresponding to the fingers, even though this leakage is substantially smaller than the one obtained by Liu and Zhang [2004] for the same model. Moreover, the palm and the back of the hand belong to different clusters since the combined metric is not enough to capture all the volumetric information. These limitations could be overcome if we include in the definition of the combined metric a part-aware distance [Liu et al. 2009].

5.3. Mesh Segmentation Quantitative Performance

In this section we study the behavior of our mesh segmentation algorithm FSS through several examples of the Princeton Segmentation Benchmark [Chen et al. 2009] for evaluation of 3D mesh segmentation algorithms. This benchmark comprises a data set with 380 surface meshes of 19 different object categories. It also provides a ground-truth corpus of 4300 human segmentation.

In the next examples we observe that the algorithm proposed in this paper, based on the computation of few columns of the affinity matrix, produces segmentations that compare to the results obtained using the full affinity matrix. Recall that it doesn't mean that the segmentation obtained with few columns of the affinity matrix W is always good, but that it is as good as the one obtained computing all columns of W . In other words, if we select carefully which columns of W are computed, then the quality of the results essentially depends on how good the selected metric reflects the features of the triangulation. In our experiments we compute the distance between triangles using several metrics: the geodesic distance, the angular distance, the product of them, and the SDF distance [Shapira et al. 2008].

Usually, the quantitative evaluation of a segmentation algorithm is done by comparing the automatic segmentation with one or more reference segmentations of the ground-truth corpus. In the literature one can find several metrics to evaluate quantita-

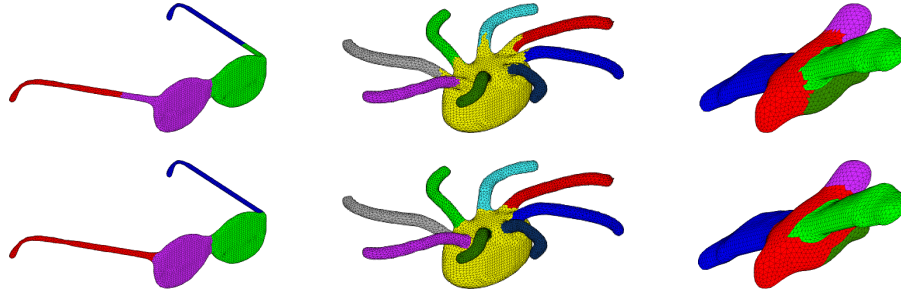


Figure 7. Top: Segmentations obtained computing few columns of the matrix W . Bottom: Ground-truth segmentations. Left: Segmentation of the sunglasses based on geodesic distance and obtained computing 2% of the 8324 columns of W . Center: Segmentation of the octopus based on the angular distance and obtained computing 1% of the 11,888 columns of W . Right: Segmentation of the bird based on the SDF distance and obtained computing 1% of the 6312 columns of W .

tively the similarity between two segmentations of a triangulated surface [Benhabiles et al. 2010; Chen et al. 2009]. In this section we employ two different non-parametric measures: the Jaccard index JI [Fowlkes and Mallows 1983] and the Rand index RI [Rand 1971]. The segmentation of a mesh with n triangles may be described by a vector $s = (s_1, \dots, s_n)$, where s_j is the index of the cluster to which the j -th triangle belongs. Given two segmentations s_a and s_b of the same mesh, we denote by $JI(s_a, s_b)$ and $RI(s_a, s_b)$ the similarity between them according to the Jaccard and Rand indexes, respectively. For both indexes it holds that $0 \leq JI(s_a, s_b) \leq 1$ and $0 \leq RI(s_a, s_b) \leq 1$, where the value 1 corresponds to the maximal similarity, i.e., $JI(s_a, s_b) = 1$ or $RI(s_a, s_b) = 1$ means that segmentations s_a and s_b are identical. In the experiments we compute Jaccard and Rand distances between s_a and s_b given by $d_J(s_a, s_b) := 1 - JI(s_a, s_b)$ and $d_R(s_a, s_b) := 1 - RI(s_a, s_b)$, respectively.

In Figure 7 we show the results obtained for three models of the Princeton Segmentation Benchmark: the sunglasses (model 42), the octopus (model 121), and the bird (model 243).

For all models, Rand and Jaccard distances are computed comparing the automatic segmentation (top row) with the ground-truth segmentation (bottom row). Table 1 shows the values of the Rand and Jaccard distances as well as the percent of columns of the affinity matrix used to obtain the automatic segmentation and the metric employed for computing the distance between triangles. In these examples a low percent of columns provides good segmentations.

In Figure 8 we illustrate that the quality of the automatic segmentation with few columns of W strongly depends on the capability of the selected metric to capture the features of the mesh. In this example we show that even if we compute the full affinity matrix W , the resulting automatic segmentation is far away from the ground-

Model	Metric	%	Distance	
			Rand	Jaccard
Sunglasses	Geodesic	2	0.062	0.161
Octopus	Angular	1	0.052	0.100
Bird	SDF	1	0.048	0.295

Table 1. Rand and Jaccard distances (d_R and d_J , respectively) between the automatic segmentation and the ground-truth segmentation.

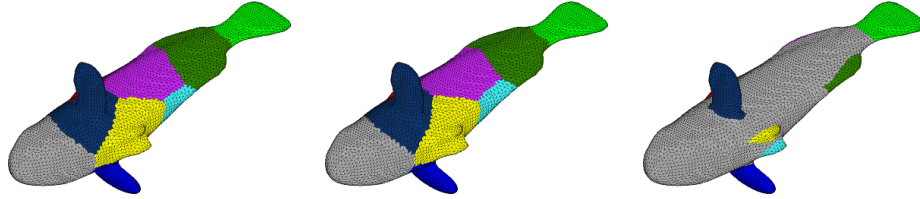


Figure 8. Three segmentations of a fish (model 225 of the Princeton Segmentation Benchmark). Left and center: Segmentations based on the geodesic distance obtained using 0.5% and 100% of the 12,148 faces. The Rand and Jaccard distances between them are $d_R = 0.010$ and $d_J = 0.064$, respectively. Right: Ground-truth segmentation. The Rand and Jaccard distances between the segmentation obtained computing all the columns of W and the ground-truth segmentation are $d_R = 0.415$ and $d_J = 0.742$, respectively.

truth segmentation, as shown by the values of d_R and d_J between these segmentations ($d_R = 0.415$, $d_J = 0.742$). Conversely, the automatic segmentations obtained computing 0.5% and 100% of the columns of W are very similar, since the values of d_R and d_J between them are small ($d_R = 0.010$, $d_J = 0.064$). Hence, the segmentation with few columns is also not good in comparison with the ground-truth segmentation. Finally, in this example we also observe that the body of the fish is subdivided into clusters with similar areas. As [Chen et al. \[2009\]](#) pointed out, this behavior is typical of segmentations based on the k -means algorithm.

5.4. Quality of Segmentations: A Different Way of Measuring

In the experiments of this section we use a different approach to measure the quality of the segmentation. Instead of comparing the automatic segmentation, obtained computing k columns of the affinity matrix, with the ground truth of the corpus [[Chen et al. 2009](#)], we compare it with the segmentation obtained using all columns of the affinity matrix. In our opinion, this comparison is fairer, since the simple metrics (geodesic, angular, and SDF distances) that we have used to compute the distance and affinity matrices are not always enough to produce good segmentations. Hence, the comparison of the automatic segmentations with ground-truth corpus segmentations does not help in the sense of proving that the results with few well-selected columns

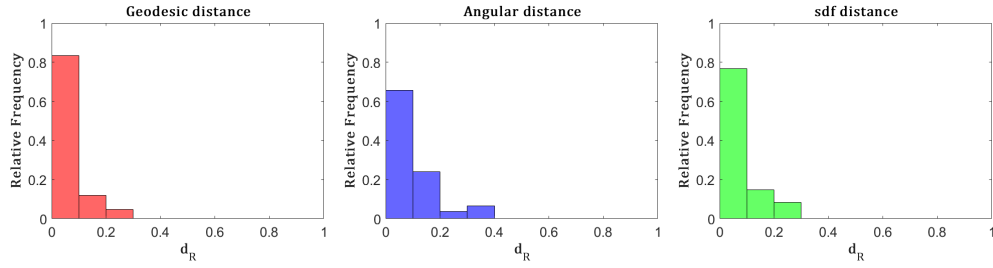


Figure 9. Relative frequency histograms of the Rand distance values, d_R , between the segmentations obtained with the full matrix W and the segmentation obtained with the same metric for k columns of W , with k equal to 0.5%, 1%, 2%, 5%, 10%, and 25% of n . From left to right: Geodesic, angular, and SDF metric.

of the affinity matrix are of quality quite similar to the results obtained computing the full affinity matrix.

Applying our segmentation method, 18 meshes of the Princeton Segmentation Benchmark [Chen et al. 2009] are segmented using k columns of the affinity matrix W , where $k = 0.5\%$, 1% , 2% , 5% , 10% and 25% of the total number n of faces. The Rand distance d_R between the segmentations obtained for k columns and the segmentation obtained with the same distance for the full matrix W is computed. Three different metrics are considered: the geodesic, angular, and SDF metrics. For each of these metrics, Figure 9 shows the histogram of the relative frequency of the Rand distance values between the segmentations obtained for k columns and the segmentation obtained with the same metric for the full matrix W . This experiment shows that, with high probability, the segmentations with a small number of columns k are very close to the corresponding segmentation obtained for the full matrix W .

5.5. The β Curve

The graph of β_k from Equation (18) as a function of k has an L shape, similar to the singular values curve, which decreases very fast for small values of k ; see Figure 4. This suggests that the β curve could be used to propose a lower bound for the size k of the sample that furnishes a projection H^k providing a good approximation to W . The top row of Figure 10 shows the curve $(k, \beta_k/\beta_1)$, for $k = 1, \dots, k_{\max}$, for three models of the Princeton Segmentation Benchmark: hand (model 185.off with $k_{\max} = 497$), bearing (model 341.off with $k_{\max} = 66$), and octopus (model 125.off with $k_{\max} = 187$). We consider different distances: geodesic for the hand model, angular for the bearing model, and SDF for the octopus model. Observe that these curves decrease very fast for small values of k and tend slowly to 0 when k goes to n .

For several values of k , Table 2 shows the Rand and Jaccard distances between the automatic segmentation and the ground-truth segmentation of the Princeton Segmentation Benchmark. The smallest number of columns for which the slope of the β

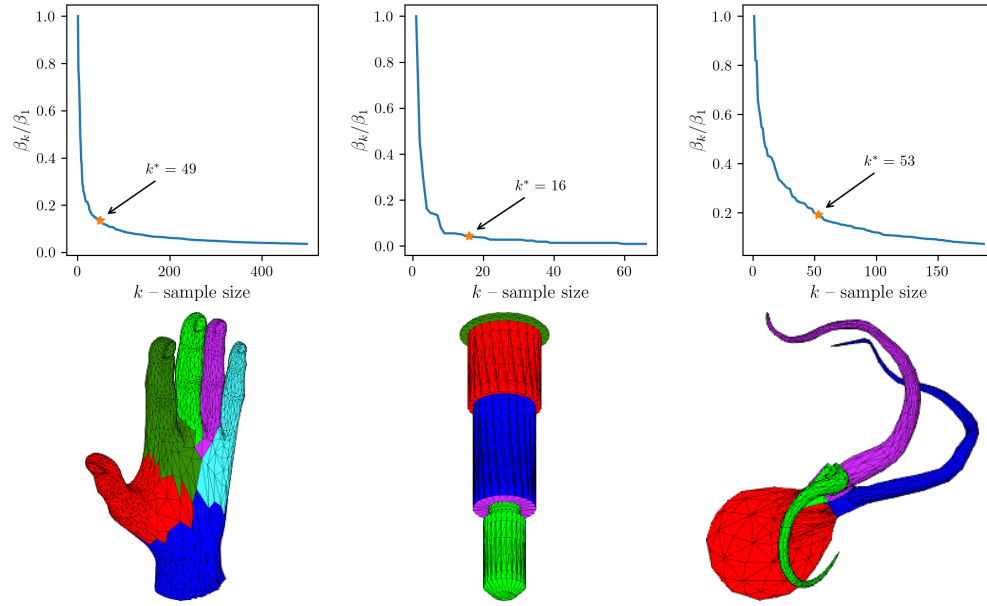


Figure 10. Top: Normalized β_k/β_1 curve. Bottom: Segmentation of the models based on W^{k^*} . Different distances are used: geodesic for the hand (left), angular for the bearing (center), and SDF for the octopus (right).

curve may be considered as very small are marked with *. Observe that for any fixed model and all considered values of k , the values of the Rand and Jaccard distances between the automatic segmentation and the ground-truth segmentation are very similar. In the bottom row of Figure 10, the segmentations corresponding to k^* are shown.

Model	Metric	k	Distance	
			Rand	Jaccard
Hand (4974 faces)	Geodesic	49*	0.124	0.303
		497	0.124	0.303
		124	0.123	0.302
Bearing (3322 faces)	Angular	16*	0.033	0.452
		66	0.033	0.452
		166	0.028	0.452
Octopus (2682 faces)	SDF	53*	0.043	0.105
		187	0.040	0.100
		321	0.040	0.099

Table 2. Rand and Jaccard distances between the proposed automatic segmentation and the ground-truth segmentation. The automatic segmentation was obtained using the metric indicated in the second column and computing the number of columns of the affinity matrix indicated in the third column.

5.6. Comparing Our Method with the Spectral Approach

Since the FSS method proposed in this paper has some points of contact with the one introduced by Liu and Zhang [2004] and later improved by Liu et al. [2006], in this section we compare the segmentations obtained with both approaches. The later one applies Nyström’s method to approximate the spectral embeddings of faces of the triangulation. To avoid the expensive computation of the normalized matrix $Q = M^{-1/2}WM^{-1/2}$, with M given by Equation (2), and its largest eigenvectors, Nyström’s method computes approximately the largest eigenvectors of Q , from a small sample of its rows (or columns) and the solution of a small-scale eigenvalue problem (see Section 3.1). The final step consists in applying k -means to the rows of the matrix of eigenvectors of Q . The selection of the sample has a strong influence on the accuracy of the eigenvectors.

No comments on the recommended relationship between the number of clusters and the number of eigenvectors are included by Liu et al. [2006]. In the numerical experiments reported here, the results with the spectral method are obtained setting the number of eigenvectors equal to the number of clusters, as it is customary in spectral segmentation [von Luxburg 2007]. Furthermore, the same sample of max-min farthest faces is used to select the columns of W to be computed by our FSS method and also for the Nyström approximation of the largest eigenvectors of W . Moreover, as suggested by Fowlkes et al. [2004], Nyström-approximated eigenvectors of Q are orthogonalized before applying k -means clustering.

Figure 11 shows the segmentations based on the SDF distance of several models, using the spectral method with Nyström approximation and using FSS. Table 3 shows the values of the Rand and Jaccard distances between the corresponding segmentations. In general we observe that even when the segmentations are different, the quality of them is similar. The Rand distances between segmentations are very small in all cases, but the Jaccard distances are larger, reflecting better the visual differences.

In our last example we show an unexpected artifact that we have observed. Sometimes the spectral segmentation obtained using Nyström approximation produces non-connected clusters. In contrast, our segmentation approach using the same sample of columns of W always produces connected clusters (recall Lemma 1). In Figure 12 (left) we show the spectral segmentation of the woman model obtained with the normalized matrix Q . The center and right images of this figure show respectively the spectral segmentation with Nyström approximation and with the method proposed in this paper, both using 5% of the columns of W corresponding to the farthest triangles in the SDF distance. The described artifact becomes evident in the spectral segmentation with Nyström approximation.

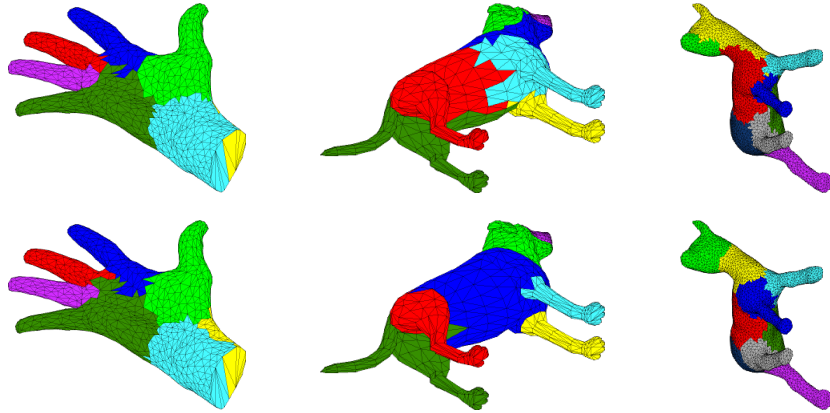


Figure 11. Segmentations based on the SDF distance. Top: Spectral segmentations with Nyström approximation; the number of eigenvectors is equal to the number of clusters. Bottom: Segmentations obtained with our FSS method. The sample of columns of the affinity matrix is the same in both approaches: 0.5% of n for the hand (left), 25% of n for dog (center), and 5% of n for fawn (right).

Model	%	Distance	
		Rand	Jaccard
Hand	0.5	0.054	0.229
Dog	25	0.082	0.406
Fawn	5	0.084	0.397

Table 3. Rand and Jaccard distances between our segmentation and the segmentation produced by the spectral approach using Nyström approximation. The sample of columns of W is the same for both methods. The second column of the table shows the size of the sample that is the indicated percent of the total number of triangles.

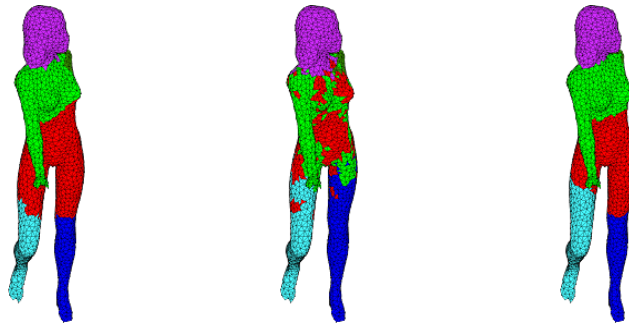


Figure 12. Left: Spectral segmentation computing all columns of the affinity matrix W . Center: Spectral segmentation with Nyström approximation computing 5% of the columns of W . Right: Our segmentation based in the same sample of columns of W .

6. Conclusions and Future Work

6.1. Conclusions

We have proposed a segmentation method for triangulated surfaces that only depends on a metric to quantify the distance between triangles and on the selection of a sample of few triangles. The proposed method computes the weighted dual graph of the triangulation with weights equal to the distances between neighboring triangles. The k farthest triangles in the chosen metric are used to compute a rectangular affinity matrix W^k of order $n \times k$, where the number k of columns is much smaller than the total number of triangles n . Rows of W^k encode the similarity between all triangles and the k triangles of the sample. Thus, clustering the rows of W^k happens to be consistent with the results of clustering the rows of the full affinity matrix W and no artifact appears. Hence, a valid strategy to solve the segmentation problem consists in clustering the rows of W^k by using, for instance, the k -means algorithm.

From the theoretical point of view, the problem of reducing the dimensionality for clustering is strongly connected with the low-rank approximation of the matrix containing the data to be clustered, which in our context is the affinity matrix W . In this sense, we have proved that for any sample of k indexes, the rank- k approximation of W obtained projecting it on the space generated by the columns of W with indexes in the sample, coincides with the rank- k approximation obtained projecting W on the space generated by its approximated eigenvectors, computed by Nyström's method with the same sample of columns of W . Moreover, it is shown that if the columns of W^k correspond to the k farthest triangles in the selected metric, then the proximity relationship among the rows of W^k tends to faithfully reflect the proximity among the corresponding rows of W .

In practice, our experiments have confirmed that this occurs even for relatively small k , resulting in a low computational cost for our method. Multiple experiments with a large variety of 3D triangular meshes were performed, and they have shown that the segmentations obtained when k is less than 10% of n are as good as those obtained from clustering the rows of the full matrix W . We have also observed that the quality of the results, objectively measured in terms of Rand and Jaccard distances between the automatic and the ground-truth segmentations, depends strongly on the capability of the selected metric of capturing the geometrical features of the mesh. Our experiments with geodesic, angular, and SDF distances show that none is enough to produce good segmentations in all cases. A combination of two or more metrics usually leads to better results.

Compared to other segmentation methods considered in the literature, the segmentation method proposed in this paper has several advantages. First, it does not depend on parameters that must be tuned by hand. Second, it is very flexible since it can handle any metric to define the distance between triangles. Finally, it is very

cheap, with a computational cost of $\mathcal{O}(knm)$, where m is the cost of computing the distance between two faces of the triangulation. In this sense, the proposed method is cheaper than spectral segmentation methods, which in the best case (when Nyström approximation is used) compute additionally the eigenvectors of an order- k matrix, with an extra cost of $\mathcal{O}((n - k)k^2) + \mathcal{O}(k^3)$ operations.

6.2. Future Work

In the present work, we intentionally focus on simple single segmentation fields on 3D meshes, and the clustering is obtained by applying a well-known clustering algorithm, k -means++, in order to achieve straightforwardly a fair comparison of our segmentation method with the spectral method. Nevertheless, the FSS method may be extended to more complicated scenarios, where several attributes are combined in a single segmentation field [Liu et al. 2009; Wang et al. 2014], or where several multi-view clustering methods have been proposed to integrate without supervision multiple information from the data [Cai et al. 2013; Huang et al. 2012]. Therefore, in the future, we plan to investigate the advantages of replacing Nyström’s method with ours, in order to propose more efficient algorithms in terms of computational and memory complexity as well as to obtain segmentations without artifacts (compare, for instance, to existing works on the subject [Elgohary et al. 2014; Yang et al. 2012]).

It is worth mentioning that the basic ideas of our FSS method—to select a small subset of farthest faces as landmarks and to compute the distances from each node in the graph to those landmarks—could be straightforwardly used in other segmentation or clustering problems, such as the segmentation of digital images [Miao and Chen 2016] or the optimization of surface quality in 3D printing [Li and Peng 2020; Wang et al. 2016]. Furthermore, the novel theoretical results and the experiments with the weighted graphs of surface triangulations presented in this work strongly suggest that FSS, complemented with our deterministic criterion to choose the sample size k , is a good candidate for the selection of scalable landmarks strategies for shortest-path computation in large networks [Potamias et al. 2009].

References

- ARTHUR, D. AND VASSILVITSKII, S. K-Means++: The advantages of careful seeding. In *Proceedings of the Eighteenth Annual ACM-SIAM Symposium on Discrete Algorithms, SODA’07*, pages 1027–1035. Society for Industrial and Applied Mathematics, 2007. URL: <https://dl.acm.org/doi/10.5555/1283383.1283494>. 158
- BENHABILES, H., VANDEBORRE, J. P., LAVOUÉ, G., AND DAOUDI, M. A comparative study of existing metrics for 3D-mesh segmentation evaluation. *The Visual Computer*, 26(12):1451–1466, 2010. URL: <https://link.springer.com/article/10.1007/s00371-010-0494-2>. 161

- BOUTSIDIS, C., ZOUZIAS, A., MAHONEY, M. W., AND DRINEAS, P. Randomized dimensionality reduction for k -Means clustering. *IEEE Transactions on Information Theory*, 61(2):1045–1062, 2015. URL: <https://doi.org/10.1109/TIT.2014.2375327>. 145
- BRAND, M. AND HUANG, K. A unifying theorem for spectral embedding and clustering. In BISHOP, C. M. AND FREY, B. J., editors, *Proceedings of the Ninth International Workshop on Artificial Intelligence and Statistics*, volume R4 of *Proceedings of Machine Learning Research*, pages 41–48. PMLR, 2003. URL: <http://proceedings.mlr.press/r4/brand03a.html>. Reissued by PMLR on April 1, 2021. 142, 146
- CAI, X., NIE, F., AND HUANG, H. Multi-view K-Means clustering on big data. In *Proceedings of the Twenty-Third International Joint Conference on Artificial Intelligence, IJCAI '13*, pages 2598–2604. AAAI Press, 2013. URL: <https://www.ijcai.org/Proceedings/13/Papers/383.pdf>. 168
- CHEN, X., GOLOVINSKIY, A., AND FUNKHOUSER, T. A benchmark for 3D mesh segmentation. *ACM Transactions on Graphics*, 28(3):73:1–73:12, 2009. URL: <https://doi.org/10.1145/1531326.1531379>. 157, 160, 161, 162, 163
- DE GOES, F., GOLDENSTEIN, S., AND VELHO, L. A hierarchical segmentation of articulated bodies. In *Proceedings of the Symposium on Geometry Processing*, SGP '08, pages 1349–1356, Goslar, DEU, 2008. Eurographics Association. URL: <https://dl.acm.org/doi/10.5555/1731309.1731315>. 141, 145
- DE SILVA, B. AND TENENBAUM, J. B. Sparse multidimensional scaling using landmark points. Technical report, Stanford University, 2004. URL: https://graphics.stanford.edu/courses/cs468-05-winter/Papers/Landmarks/Silva_landmarks5.pdf. 147
- ECKART, C. AND YOUNG, G. The approximation of one matrix by another of lower rank. *Psychometrika*, 1(3):211–218, 1936. URL: <https://doi.org/10.1007/BF02288367>. 146
- ELAD, A. AND KIMMEL, R. On bending invariant signatures for surfaces. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 25(10):1285–1295, 2003. URL: <https://doi.org/10.1109/TPAMI.2003.1233902>. 146
- ELGOHARY, A., FARAHAT, A. K., KAMEL, M. S., AND KARRAY, F. Embed and conquer: Scalable embeddings for kernel k -Means on MapReduce. In *Proceedings of the 2014 SIAM International Conference on Data Mining (SDM)*, pages 425–433. Society for Industrial and Applied Mathematics, 2014. URL: <https://doi.org/10.1137/1.9781611973440.49>. 168
- FOWLKES, C., BELONGIE, S., CHUNG, F., AND MALIK, J. Spectral grouping using the Nyström method. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 26(2):214–225, 2004. URL: <https://doi.org/10.1109/TPAMI.2004.1262185>. 147, 165
- FOWLKES, E. B. AND MALLOWS, C. L. A method for comparing two hierarchical clusterings. *Journal of the American Statistical Association*, 78(383):553–569, 1983. URL: <https://doi.org/10.2307/2288117>. 161

- GAL, R., SHAMIR, A., AND COHEN-OR, D. Pose-oblivious shape signature. *IEEE Transactions on Visualization and Computer Graphics*, 13(2):261–271, 2007. URL: <https://doi.org/10.1109/TVCG.2007.45>. 143
- GOTSMAN, C. On graph partitioning, spectral analysis, and digital mesh processing. In *2003 Shape Modeling International*, pages 165–171. IEEE, 2003. URL: <https://doi.org/10.1109/SMI.2003.1199613>. 146
- HOLODNAK, J. T., IPSEN, I. C. F., AND WENTWORTH, T. Conditioning of leverage scores and computation by QR decomposition. *SIAM Journal on Matrix Analysis and Application*, 36(3):1143–1163, 2015. URL: <https://doi.org/10.1137/140988541>. 148
- HUANG, H. C., CHUANG, Y. Y., AND CHEN, C. S. Affinity aggregation for spectral clustering. In *2012 IEEE Conference on Computer Vision and Pattern Recognition*, pages 773–780. IEEE, 2012. URL: <https://doi.org/10.1109/CVPR.2012.6247748>. 168
- JAIN, V., ZHANG, H., AND VAN KAICK, O. Non-rigid spectral correspondence of triangle meshes. *International Journal of Shape Modeling*, 13(1):101–124, 2007. URL: <https://doi.org/10.1142/S0218654307000968>. 146
- JOLLIFFE, I. T. Discarding variables in a principal component analysis. I: Artificial data. *Journal of the Royal Statistical Society, Series C (Applied Statistics)*, 21(2): 160–173, 1972. URL: <https://academic.oup.com/jrssc/article/21/2/160/6882622>. 148
- KAMOUSHI, P., LAZARD, S., MAHESHWARI, A., AND WUHRER, S. Analysis of farthest point sampling for approximating geodesics in a graph. *Computational Geometry*, 57:1–7, 2016. URL: <https://doi.org/10.1016/j.comgeo.2016.05.005>. 149
- KATZ, S. AND TAL, A. Hierarchical mesh decomposition using fuzzy clustering and cuts. *ACM Transactions on Graphics*, 22(3):954–961, 2003. URL: <https://doi.org/10.1145/882262.882369>. 141, 143, 144
- KATZ, S., LEIFMAN, G., AND TAL, A. Mesh segmentation using feature point and core extraction. *The Visual Computer*, 21:649–658, 2005. URL: <https://doi.org/10.1007/s00371-005-0344-9>. 142
- KLEINBERG, J., SLIVKINS, A., AND WEXLER, T. Triangulation and embedding using small sets of beacons. In *45th Annual IEEE Symposium on Foundations of Computer Science*, pages 444–453. IEEE, 2004. URL: <https://doi.org/10.1109/FOCS.2004.70>. 149
- KOSCHAN, A. F. Perception-based 3D triangle mesh segmentation using fast marching watersheds. In *2003 IEEE Computer Society Conference on Computer Vision and Pattern Recognition, 2003, Proceedings*, volume 2, page II. IEEE, 2003. URL: <https://doi.org/10.1109/CVPR.2003.1211448>. 141
- KUMAR, S., MOHRI, M., AND TALWALKAR, A. Sampling methods for the Nyström method. *Journal of Machine Learning Research*, 13:981–1006, 2012. URL: <https://www.jmlr.org/papers/v13/kumar12a.html>. 147

- LEE, Y., LEE, S., SHAMIR, A., COHEN-OR, D., AND SEIDEL, H. P. Mesh scissoring with minima rule and part salience. *Computer Aided Geometric Design*, 22(5):444–465, 2005. URL: <https://dl.acm.org/doi/abs/10.5555/1099154.1648425>. 141
- LI, R. AND PENG, Q. Surface quality improvement and support material reduction in 3D printed shell products based on efficient spectral clustering. *The International Journal of Advanced Manufacturing Technology*, 107(9):4273–4286, 2020. URL: <https://doi.org/10.1007/s00170-020-05299-6>. 141, 168
- LIU, R. AND ZHANG, H. Segmentation of 3D meshes through spectral clustering. In *Proceedings of the Computer Graphics and Applications, 12th Pacific Conference*, PG '04, pages 298–305. IEEE Computer Society, 2004. URL: <https://doi.org/10.1109/PCCGA.2004.1348360>. 141, 142, 145, 146, 155, 156, 160, 165
- LIU, R. AND ZHANG, H. Mesh segmentation via spectral embedding and contour analysis. *Computer Graphics Forum*, 26(3):385–394, 2007. URL: <https://doi.org/10.1111/j.1467-8659.2007.01061.x>. 142, 145
- LIU, R., JAIN, V., AND ZHANG, H. Sub-sampling for efficient spectral mesh processing. In NISHITA, T., PENG, Q., AND SEIDEL, H.-P., editors, *Advances in Computer Graphics*, volume 4035 of *Lecture Notes in Computer Science*, pages 172–184, Berlin, 2006. Springer. URL: https://link.springer.com/chapter/10.1007/11784203_15. 142, 145, 146, 147, 156, 165
- LIU, R., ZHANG, H., SHAMIR, A., AND COHEN-OR, D. A part-aware surface metric for shape analysis. *Computer Graphics Forum (Special Issue of Eurographics)*, 28(2):397–406, 2009. URL: <https://doi.org/10.1111/j.1467-8659.2009.01379.x>. 141, 143, 144, 160, 168
- LLOYD, S. Least squares quantization in PCM. *IEEE Transactions on Information Theory*, 28(2):129–137, 1982. URL: <https://doi.org/10.1109/TIT.1982.1056489>. 145, 156
- MAHONEY, M. W. Randomized algorithms for matrices and data. *Foundations and Trends in Machine Learning*, 3(2):123–224, 2011. URL: <https://doi.org/10.1561/22000000035>. 147, 148, 149
- MIAO, J. AND CHEN, D. A spectral clustering image segmentation algorithm based on Nyström approximation. *DEStech Transactions on Computer Science and Engineering*, 2016. URL: <https://doi.org/10.12783/dtcse/cmsam2016/3596>. 168
- MULLEN, P., TONG, Y., ALLIEZ, P., AND DESBRUN, M. Spectral conformal parameterization. In *Proceedings of the Symposium on Geometry Processing*, SGP '08, pages 1487–1494, Goslar, DEU, 2008. Eurographics Association. URL: <https://dl.acm.org/doi/10.5555/1731309.1731335>. 146
- NG, A. Y., JORDAN, M. I., AND WEISS, Y. On spectral clustering: Analysis and an algorithm. In *Proceedings of the 14th International Conference on Neural Information Processing Systems: Natural and Synthetic*, NIPS'01, pages 849–856, Cambridge, MA, 2001. MIT Press. URL: https://proceedings.neurips.cc/paper_files/paper/2001/hash/801272ee79cfde7fa5960571fee36b9b-Abstract.html. 145

- POTAMIAS, M., BONCHI, F., CASTILLO, C., AND GIONIS, A. Fast shortest path distance estimation in large networks. In *Proceedings of the 18th ACM Conference on Information and Knowledge Management, CIKM '09*, pages 867–876. Association for Computing Machinery, 2009. URL: <https://doi.org/10.1145/1645953.1646063>. 149, 168
- RAND, W. M. Objective criteria for the evaluation of clustering methods. *Journal of the American Statistical Association*, 66(336):846–850, 1971. URL: <https://doi.org/10.2307/2284239>. 161
- REUTER, M., BIASOTTI, S., GIORGI, D., PATANÈ, G., AND SPAGNUOLO, M. Technical section: Discrete Laplace-Beltrami operators for shape analysis and segmentation. *Computers & Graphics*, 33(3):381–390, 2009. URL: <https://doi.org/10.1016/j.cag.2009.03.005>. 146
- SHAMIR, A. A survey on mesh segmentation techniques. *Computer Graphics Forum*, 27(6):1539–1556, 2008. URL: <https://doi.org/10.1111/j.1467-8659.2007.01103.x>. 141
- SHAPIRA, L., SHAMIR, A., AND COHEN-OR, D. Consistent mesh partitioning and skeletonisation using the shape diameter function. *The Visual Computer*, 24(4):249–258, 2008. URL: <https://doi.org/10.1007/s00371-007-0197-5>. 143, 144, 156, 160
- SHI, J. AND MALIK, J. Normalized cuts and image segmentation. In *Proceedings of the 1997 Conference on Computer Vision and Pattern Recognition (CVPR '97)*, CVPR '97, pages 731–737. IEEE Computer Society, 1997. URL: <https://doi.org/10.1109/CVPR.1997.609407>. 145
- VON LUXBURG, U. A tutorial on spectral clustering. *Statistics and Computing*, 17(4):395–416, 2007. URL: <https://doi.org/10.1007/s11222-007-9033-z>. 145, 165
- WANG, H., LU, T., AU, O. K. C., AND TAI, C. L. Spectral 3D mesh segmentation with a novel single segmentation field. *Graphical Models (Geometric Modeling and Processing 2014)*, 76(5):440–456, 2014. URL: <https://doi.org/10.1016/j.gmod.2014.04.009>. 143, 156, 168
- WANG, W. M., ZANNI, C., AND KOBELT, L. Improved surface quality in 3D printing by optimizing the printing direction. In *Proceedings of the 37th Annual Conference of the European Association for Computer Graphics, EG '16*, pages 59–70, Goslar, DEU, 2016. Eurographics Association. URL: <https://doi.org/10.1111/cgfm.12811>. 141, 168
- YANG, T., LI, Y.-F., MAHDAVI, M., JIN, R., AND ZHOU, Z.-H. Nyström method vs random fourier features: a theoretical and empirical comparison. In *Proceedings of the 26th International Conference on Neural Information Processing Systems - Volume 1, NIPS'12*, pages 476–484, Red Hook, NY, 2012. Curran Associates Inc. URL: https://proceedings.neurips.cc/paper_files/paper/2012/hash/621bf66ddb7c962aa0d22ac97d69b793-Abstract.html. 168
- ZHANG, H. AND LIU, R. Mesh segmentation via recursive and visually salient spectral cuts. In *Vision, Modeling, and Visualization 2005: Proceedings, November 16-18, 2005*,

Erlangen, Germany, pages 429–436. IOS Press, 2005. URL: https://www.cs.sfu.ca/~haoz/pubs/zhang_liu_vmv05.pdf. 141, 142

ZHANG, H., VAN KAICK, O., AND DYER, R. Spectral mesh processing. *Computer Graphics Forum*, 29(6):1865–1894, 2010. URL: <https://doi.org/10.1111/j.1467-8659.2010.01655.x>. 142

Author Contact Information

Victoria Hernández Mederos
Instituto de Cibernética, Matemática y Física
La Habana, Cuba

Dimas Martínez
Departamento de Matemática
Universidade Federal do Amazonas
Manaus, Brazil

Jorge Estrada Sarlabous
Instituto de Cibernética, Matemática y Física
La Habana, Cuba

Valia Guerra Ones
Departamento de Análisis Matemático
Universidad de La Laguna
Tenerife, Spain

V. Hernández Mederos, D. Martínez, J. Estrada Sarlabous, and V. Guerra Ones, Farthest Sampling Segmentation of Triangulated Surfaces, *Journal of Computer Graphics Techniques (JCGT)*, vol. 14, no. 1, 140–173, 2025
<http://jcgt.org/published/0014/01/07/>

Received: 2022-10-05

Recommended: 2025-02-01

Published: 2025-04-02

Corresponding Editor: David Eberly

Editor-in-Chief: Eric Haines

© 2025 V. Hernández Mederos, D. Martínez, J. Estrada Sarlabous, and V. Guerra Ones (the Authors).

The Authors provide this document (the Work) under the Creative Commons CC BY-ND 3.0 license available online at <http://creativecommons.org/licenses/by-nd/3.0/>. The Authors further grant permission for reuse of images and text from the first page of the Work, provided that the reuse is for the purpose of promoting and/or summarizing the Work in scholarly venues and that any reuse is accompanied by a scientific citation to the Work.

