

Manifold Learning, Topological Data Analysis, and Their Applications to Medical Imaging

By

Kun Meng

Sc.M., Brown University, 2018

M.S., Beijing Normal University, 2015

B.S., Tianjin Polytechnic University, 2012

A dissertation submitted in partial fulfillment of the requirements for
the Degree of Doctor of Philosophy
in the Department of Biostatistics at Brown University

PROVIDENCE, RHODE ISLAND

April 2022

©Copyright 2022 by Kun Meng

This dissertation by Kun Meng is accepted in its present form
by the Department of Biostatistics as satisfying the dissertation requirement for
the degree of Doctor of Philosophy.

Date

Ani Eloyan, Advisor

Recommended to the Graduate Council

Date

Constantine Gatsonis, Reader

Date

Matthew Harrison, Reader

Date

Lorin Crawford, Reader

Approved by the Graduate Council

Date

Andrew Campbell, Dean of the Graduate School

Acknowledgements

First and foremost, I would like to thank my thesis advisor Dr. Ani Eloyan for her mentorship and friendship. When I came to the Department of Biostatistics at Brown University in 2016, I had very limited knowledge of statistics, had never really read any papers on statistics, and was almost not able to speak English. Ani was very patient and started teaching me every skill in statistical research. After my six years at Brown — two years as a master's student and four years as a Ph.D. student, I have now finished my dissertation covering three different topics in statistics, published a paper in a top statistical journal, and got the offer for my first job in academia — Prager assistant professorship in the Division of Applied Mathematics at Brown University. All these achievements are attributable to Ani.

In addition to my thesis advisor, I want to thank other thesis committee members. I would like to thank Dr. Constantine Gatsonis. Constantine taught me the Statistical Inference II course — the first building block of the bridge in my mind connecting pure mathematics and learning theory; he also gave me many suggestions on my research and career path. I would like to thank Dr. Lorin Crawford. The last chapter of my dissertation depends on one part of Lorin's dissertation; without his help, I would not even be able to start the last chapter of my dissertation herein. I would like to thank Dr. Matthew Harrison. In the last four years, I took or audited several courses taught by Matt, especially APMA 2640 — Probability Theory II. APMA 2640 was the first graduate-level probability theory course I took in my life, which provided me with the necessary mathematical tools for this dissertation. Additionally, I feel honored that I will be a colleague of Matt this summer in the Division of Applied Mathematics.

I want to thank all the teachers who taught me knowledge. I would like to particularly thank Prof. Yonghong Yao and Prof. Junqing Wang at Tianjin Polytechnic University. They taught me Mathematical Analysis and Linear Algebra courses when I was a freshman. Most of my mathematical capacity is attributable to their teaching and training.

I want to thank all my friends in both China and the US. I would like to thank Jerson Cochancela and his wife, Angela Pizzuto. Jerson, Angela, and I came to Brown as master's students in the same year, and then we got accepted to the Ph.D. programs at Brown in the same year. In the past six years, Jerson and Angela treated me as their family member and helped me greatly in every aspect of my life at Brown. I went to their wedding, and they came to my dissertation defense. I sincerely appreciate their friendship, which is definitely a lifelong one. I would like to thank Dr. Zixiong Wang and Dr. Yijian Chuan, particularly for their help and support from 2015 to 2016 — one of the darkest years in my life.

Finally, I am grateful to my family.

Summary

This dissertation aims to develop several statistical methods for learning the manifold and topology structures of data, provide the corresponding theoretical foundations, and apply the proposed methods to medical imaging. Specifically, this dissertation is concerned with manifold learning, brain networks, and topological data analysis.

The concepts of manifolds and topology were proposed in pure mathematics to investigate geometry from local and global viewpoints, respectively. With the development of mathematical statistics and the advances in statistical applications to high-dimensional data, medical imaging analysis, and neural science, the concepts of manifolds and topology gradually entered — perhaps are still entering — statistics.

Manifold learning refers to a collection of methods detecting low-dimensional and potentially nonlinear geometric structures from high-dimensional data. The detected geometric structures are required to be continuous and hopefully have some smoothness. Topological data analysis is a collection of statistical methods that find structures of data using algebraic topology. One of the simplest structures in algebraic topology is simplicial complexes. One-dimensional simplicial complexes are graphs. Brain networks are modeled as graphs, and estimating brain networks can be boiled down to learning the topology in brain science data. Therefore, the theme of this dissertation may also be “manifold and topology learning.”

In Chapter 1 — *Principal Manifold Estimation via Model Complexity Selection*, we propose a framework of principal manifolds to model high-dimensional data. This framework is based on Sobolev spaces and designed to model data of any intrinsic dimension. It includes principal component analysis and principal curve algorithm as special cases. A Sobolev embedding theorem guarantees the regularity of the proposed principal manifolds, and the theory of reproducing kernel Hilbert spaces provides an approach to the estimation of principal manifolds. We propose a novel method for model complexity selection to avoid overfitting, eliminate the effects of outliers, and improve the computation speed. Additionally, we propose a method for identifying the interiors of circle-like curves and cylinder/ball-like surfaces. The proposed approach is compared to existing methods by simulations and applied to estimate tumor surfaces and interiors in a lung cancer study.

In Chapter 2 — *Population-level Task-evoked Functional Connectivity*, we estimate brain networks using functional magnetic resonance imaging (fMRI) data. FMRI is a non-invasive and in-vivo imaging technique essential for measuring brain activity. Functional connectivity is used to study associations between brain regions either at rest or while study subjects perform tasks. In this paper, we propose a rigorous definition of task-evoked functional connectivity at the population level (ptFC). Importantly, our proposed ptFC is interpretable in the context of task-fMRI studies. An algorithm for estimating ptFC is provided. We present the performance of the proposed algorithm compared to existing functional connectivity estimation approaches using simulations. Lastly, we apply the proposed framework to estimate task-evoked functional connectivity in a motor-task study from the Human Connectome Project. We show that the proposed algorithm identifies associations regions of the brain related to the performance of motor tasks as expected.

In Chapter 3 — *Randomness and Statistical Inference of Shapes via the Smooth Euler Characteristic Transform*, we provide the mathematical foundations for the randomness of shapes and the distributions of smooth Euler characteristic transform. Based on these foundations, we propose an approach for testing hypotheses on random shapes. Simulation studies are provided to support our mathematical derivations and show the performance of our proposed hypothesis testing framework. Our discussions connect the following fields: algebraic and computational topology, probability theory and stochastic processes, Sobolev spaces and functional analysis, statistical inference, and medical imaging.

One potential future research direction is “global manifold learning” — a collection of methods combining manifold learning and topological data analysis. The Chern-Gauss-Bonnet formula $\chi(M) = \frac{1}{(2\pi)^n} \int_M \text{Pf}(\Omega)$ (see Chern (1944)) illustrates the spirit of this research direction. The Euler characteristic $\chi(M)$ of manifold M on the left-hand side of the formula is an estimand in topological data analysis; the curvature form Ω on the right-hand side of the formula presents the local structures of manifold M , which can be estimated using manifold learning. Therefore, it is natural and automatic to intertwine manifold learning with topological data analysis and develop global manifold learning. Most of the existing manifold learning approaches, including the one in Chapter 1 of this dissertation, essentially investigate local coordinate charts of differential manifolds. Parallel with the evolution of differential geometry in the twentieth century — going from local differential geometry to global differential geometry, the manifold learning theory in statistics should also go in this direction.

Contents

1	Principal Manifold Estimation via Model Complexity Selection	9
1.1	Introduction	9
1.2	Manifolds and Projection Indices	12
1.3	Principal Manifolds	14
1.3.1	Sobolev Spaces	14
1.3.2	Definition of Principal Manifolds	15
1.3.3	Two Special Cases	17
1.4	A New Approach to Principal Manifold Estimation	18
1.5	Step 1: Data Reduction	19
1.5.1	Motivation and Estimation of Mixture Density Parameters	19
1.5.2	Estimation of the Number of Mixture Components	21
1.5.3	The High-Dimensional Mixture Density Estimation (HDMDE) Algorithm	23
1.6	Step 2: Fitting, Step 3: Tuning, Model Complexity Selection	25
1.7	Simulations	27
1.8	Interior Identification	30
1.9	Analysis of Lung Cancer Tumor Data	33
1.10	Conclusions and Discussion	35
1.11	Appendix	36
2	Population-level Task-evoked Functional Connectivity	39
2.1	Introduction	39
2.2	Population-level Task-evoked Functional Connectivity (ptFC)	42
2.2.1	Definition of ptFC	43
2.2.2	Advantages of ptFC	45
2.3	ptFC Estimation (ptFCE) Algorithm	46
2.3.1	The Estimation of Reference Signals	48

2.3.2	The Estimator for $C_{kl}(\xi)$ and The ptFCE Algorithm	50
2.4	Simulations	51
2.5	Analysis of HCP Motor Data	54
2.6	Conclusions	57
2.A	Relationship between the proposed BOLD signal model and some existing models	59
2.B	Lemmas, Theorems, and Their Proofs	60
2.C	The Identifiability of Task-evoked Terms	65
2.D	The Performance of ptFCE in Estimating ptFCs	66
2.E	Data-generating Mechanisms for Simulations	68
2.E.1	Mechanism 0	70
2.E.2	Mechanism 1	73
2.E.3	Mechanism 2	73
2.F	Limitations of the Pearson correlation approach	74
2.G	Details of Beta-series Regression and Coherence Analysis	74
2.H	Future Research	75
3	Randomness and Statistical Inference of Shapes via the Smooth Euler Characteristic Transform	77
3.1	Introduction	77
3.1.1	Overview of Topological Data Analysis	78
3.1.2	Overview of Contributions	79
3.1.3	Relevant Notation and Paper Organization	80
3.2	The Smooth Euler Characteristic Transform of Shapes	82
3.2.1	The Definition of Smooth Euler Characteristic Transform	82
3.2.2	Proof-of-Concept Simulation Examples I: Deterministic Shapes	85
3.2.3	The Primitive Euler Characteristic Transform of Shapes	86
3.3	Random Shapes and Probabilistic Distributions over the SECT	88
3.3.1	Formulation of the SECT as a Random Variable	88
3.3.2	Mean and Covariance of the SECT	89
3.3.3	Proof-of-Concept Simulation Examples II: Random Shapes	91
3.4	Testing Hypotheses on Shapes	93
3.4.1	Karhunen–Loève Expansion	94
3.4.2	The Theoretical Foundation for Hypothesis Testing	95
3.4.3	The Numerical foundation for Hypothesis Testing	98
3.5	Experiments Using Simulations	100

3.6	Conclusions and Discussions	103
3.A	Overview of Persistent Diagrams	105
3.B	Computation of SECT	107
3.C	Proofs	109

Chapter 1

Principal Manifold Estimation via Model Complexity Selection

1.1 Introduction

Manifold learning is a method for modeling high-dimensional data, assuming that data are from a low-dimensional manifold and corrupted by high-dimensional noise. The dimension of the low-dimensional manifold is called the *intrinsic dimension* of data. There are two primary components of manifold learning: (i) *parameterization* - uncovering a low-dimensional description of high-dimensional data; (ii) *embedding* - finding a map relating the low-dimensional description and high-dimensional data. The two components are entangled with each other. Based on a given parameterization, embedding becomes a statistical fitting problem. In turn, projecting data to the image of an embedding map results in a parameterization (e.g., [Yue et al. \(2016\)](#)). In this paper, we propose a framework and estimation approach combining these two components. Specifically, our proposed approach constructs an embedding map from a “partial” parameterization and obtains a full parameterization from this embedding map. We define principal manifolds as minima of a functional equipped with a regularity penalty term derived as a semi-

-
- A version of this chapter appeared in [Meng and Eloyan \(2021b\)](#).
 - **Author contributions:** Kun Meng developed the theoretical framework of this chapter, proved the theorems, and conducted the simulation studies and the tumor data application. Ani Eloyan supervised this chapter. Kun Meng and Ani Eloyan wrote the version [Meng and Eloyan \(2021b\)](#). We — Kun Meng and Ani Eloyan — want to thank Dr. Matthew T. Harrison for his comments that improved the readability of this chapter and Dr. Chen Yue for providing the R function of the PS algorithm. Last but not least, we would like to thank the editorial team and anonymous referees of the *Journal of the Royal Statistical Society, Series B (Statistical Methodology)*, for their insightful comments that significantly improved the quality of this chapter.
 - **Abbreviations:** CT, computed tomography; HS, Hastie and Stuetzle (1989); iid, identically and independently distributed; KDE, kernel density estimate; KLD, Kullback-Leibler divergence; MSD, mean squared distance; PME, principal manifold estimate; PS, principal surface estimate by Yue et al. (2016); PCA, principal component analysis; PDF, probability density function; SD, standard deviation.

norm on a Sobolev space. A Sobolev embedding theorem implies the differentiability of our proposed manifolds. The novel framework of principal manifolds allows the intrinsic dimension of data to be any positive integer. The linear principal component analysis (PCA, [Jolliffe \(2002\)](#)) and principal curve algorithm ([Hastie and Stuetzle \(1989\)](#)) are special cases of this framework. We provide topological and functional analysis arguments giving mathematical foundations of our proposed principal manifold framework. To avoid overfitting and preserve the curvatures of underlying manifolds, we propose a model complexity selection method. Additionally, this method drastically reduces the computational cost and eliminates the effects of outliers. Based on this method and the theory of reproducing kernel Hilbert spaces, we propose an algorithm to estimate principal manifolds efficiently. Additionally, motivated by a problem in radiation therapy for lung cancer patients, we propose a method for identifying interiors of circle-like curves and cylinder/ball-like surfaces.

Throughout this paper, we use the following notations: (i) d and D , with $d < D$, denote the dimensions of intrinsic manifolds and the spaces into which these manifolds are embedded, respectively. (ii) $\|x\|_{\mathbb{R}^q} = (\sum_{k=1}^q x_k^2)^{1/2}$ for all $x \in \mathbb{R}^q$ and all positive integers q . (iii) Let $q_1, q_2 \in \{d, D\}$, $k \in \{1, 2, \dots, \infty\}$, and I be a subset of $\mathbb{R}^{q_1}$, $C^k(I \rightarrow \mathbb{R}^{q_2})$ denotes the collection of $I \rightarrow \mathbb{R}^{q_2}$ maps whose components have up to k^{th} continuous classical derivatives. For simplicity, $C^k(I) = C^k(I \rightarrow \mathbb{R}^1)$ and $C = C^0$. (iv) δ_x is the point mass at x (see Section 6.9 of [Rudin \(1991\)](#)). (v) L^p and $\|\cdot\|_{L^p}$, $p \in [1, \infty]$, denote *Lebesgue spaces* and their norms (see Chapter 2 of [Adams and Fournier \(2003\)](#)).

A considerable amount of work has been done for parameterization and embedding tasks. ISOMAP ([Tenenbaum et al. \(2000\)](#)), locally linear embedding ([Roweis and Saul \(2000\)](#)), and Laplacian eigenmaps ([Belkin and Niyogi \(2003\)](#)) constructed parameterizations of high-dimensional data. [Hastie and Stuetzle \(1989\)](#) (hereafter HS) proposed a *principal curve* framework and algorithm for the embedding task. HS defined principal curves as follows.

Definition 1.1.1. (Part I) Let $I \subset \mathbb{R}^1$ be a closed and possibly infinite interval. Suppose a map $f : I \rightarrow \mathbb{R}^D$ satisfies the conditions (referred to as HS conditions throughout this paper): (i) $f \in C^\infty(I \rightarrow \mathbb{R}^D)$; (ii) $\|f'(t)\|_{\mathbb{R}^D} = 1$ for all $t \in I$, i.e., f is arc-length parameterized; (iii) f does not self intersect, i.e. $t_1 \neq t_2$ implies $f(t_1) \neq f(t_2)$; (iv) $\int_{\{t: f(t) \in B\}} dt < \infty$ for any finite ball B in $\mathbb{R}^D$. Then $\pi_f : \mathbb{R}^D \rightarrow I$ is defined as follows and called the projection index with respect to f .

$$\pi_f(x) = \sup \left\{ t \in I : \|x - f(t)\|_{\mathbb{R}^D} = \inf_{t' \in I} \|x - f(t')\|_{\mathbb{R}^D} \right\}, \quad \text{for all } x \in \mathbb{R}^D. \quad (1.1.1)$$

(Part II) Suppose X is a continuous random D -vector with finite second moments. Principal curves of X are all maps $f : I \rightarrow \mathbb{R}^D$ satisfying HS conditions and the self-consistency defined as

$$\mathbb{E}(X | \pi_f(X) = t) = f(t). \quad (1.1.2)$$

The projection index π_f is well-defined under HS conditions. However, HS conditions are restrictive due to the

following reasons: 1) condition (ii) requires principal curves to be arc-length parameterized, while the arc-length parameterization is not generalizable to higher dimensions; 2) condition (iii) rules out many curves in applications, e.g., a handwritten “8” in handwriting recognition; 3) condition (iv) is not straightforward to verify. Furthermore, the HS principal curve framework has a model bias (see Section 6 of [Hastie and Stuetzle \(1989\)](#)).

To remove the model bias in the HS principal curve framework, [Tibshirani \(1992\)](#) proposed a new principal curve framework based on a mixture model and self-consistency (1.1.2). HS showed that curves satisfying (1.1.2) are critical points of the *mean squared distance* (MSD) functional $\mathcal{D}_X(f) = \mathbb{E} \|X - f(\pi_f(X))\|_{\mathbb{R}^D}^2$. However, [Duchamp et al. \(1996\)](#) showed that these critical points may be *saddle* points, i.e., there may exist adjacent curves with smaller MSD than that of curves satisfying (1.1.2). This *saddle issue* was a flaw of the frameworks based on (1.1.2). [Gerber and Whitaker \(2013\)](#) explained the saddle issue from the “orthogonal/along” variation trade-off viewpoint and discussed the challenges stemming from this issue in model complexity selection. To remove the saddle issue, [Gerber and Whitaker \(2013\)](#) avoided using MSD $\mathcal{D}_X(f)$ and proposed a new functional $\mathcal{Q}_X(\pi) = \mathbb{E} \left\{ [\mathbb{E}(X|\pi(X)) - X]^T \frac{d}{dt} \Big|_{t=\pi(X)} \mathbb{E}(X|\pi(X) = t) \right\}$ modeling the parameterization maps $\pi : \mathbb{R}^D \rightarrow I$. This functional penalizes the non-orthogonality between fitting error $\mathbb{E}(X|\pi(X)) - X$ and curve tangent $\frac{d}{dt} \Big|_{t=\pi(X)} \mathbb{E}(X|\pi(X) = t)$. Principal curves, which satisfy (1.1.2), correspond to the minima of $\mathcal{Q}_X(\pi)$. However, the use of $\mathcal{D}_X(f)$ for measuring the discrepancy between data X and fitted $f(\pi_f(X))$ is of interest due to the interpretability of $\mathcal{D}_X(f)$. Another approach to removing the saddle issue, while using $\mathcal{D}_X(f)$, is to avoid self-consistency and define principal curves by minimizing MSD with a length constraint or a regularity penalty. [Kégl et al. \(2000\)](#) defined principal curves as the minima $\arg \min_f \{ \mathcal{D}_X(f) : f \in BV([a, b]), V_a^b(f) \leq L \}$, where $L > 0$ is pre-defined and $BV([a, b])$ is the collection of functions f on $[a, b]$ with finite total variation $V_a^b(f) < \infty$. However, functions in $BV([a, b])$ are not necessarily everywhere differentiable. Indeed, the algorithm proposed by [Kégl et al. \(2000\)](#) fits data by polygonal lines, which are only piecewise differentiable. In many applications, we expect globally differentiable curves. Additionally, the [Kégl et al. \(2000\)](#) framework is only defined for curves, i.e., the intrinsic dimension $d = 1$. [Smola et al. \(2001\)](#) proposed the *regularized principal manifold* framework, where regularized estimates of principal manifolds are minimizers of the following form.

$$\arg \min_{f \in \mathcal{F}} \left\{ \mathbb{E} \|X - f(\pi_f(X))\|_{\mathbb{R}^D}^2 + \lambda \|Pf\|_{\mathcal{H}}^2 \right\}, \quad (1.1.3)$$

where $\mathcal{F}$ is a collection of functions and P is an operator mapping f into an inner product space $\mathcal{H}$. [Smola et al. \(2001\)](#) (Example 7) showed that the [Kégl et al. \(2000\)](#) definition of principal curves is essentially the special case of (1.1.3), where $P = \frac{d}{dt}$ and $\mathcal{H} = L^2$, i.e., the penalty term in (1.1.3) is $\|f'\|_{L^2}^2 = \int_a^b \|f'\|_{\mathbb{R}}^2 dt$ (the derivative f' is defined only almost everywhere with respect to the Lebesgue measure, rather than exactly everywhere). However, the regularized principal manifold approach defined by [Smola et al. \(2001\)](#) has several limitations. The problem of selection of the tuning parameter λ , and more generally, model complexity to avoid overfitting and preserve

intrinsic curvatures remains unresolved, as well as a definition of the projecting operator P that would correspond to the tuning parameter selection. HS shows that π_f is well-defined under the restrictive HS conditions and when the intrinsic dimension $d = 1$, however, there is no discussion of assumptions on the collection $\mathcal{F}$ by Smola et al. (2001) such that the resulting π_f is well-defined. The function spaces $\mathcal{H}$ and $\mathcal{F}$ compatible with P and π_f are not defined. Our proposed principal manifold estimation approach addresses these limitations.

The paper is organized as follows. In Section 1.2, we propose a condition for defining the projection indices π_f associated with maps $f : \mathbb{R}^d \rightarrow \mathbb{R}^D$, where d is allowed to be any positive integer. This proposed condition is less restrictive and easier to verify than HS conditions. In Section 1.3, based on this condition and function spaces of Sobolev type, we define principal manifolds by minimizing MSD equipped with a total squared curvature penalty. This definition solves the differentiability problem in Kégl et al. (2000). In Section 1.4, we present the outline of our proposed *principal manifold estimation* algorithm by briefly describing its main steps. Details of these steps are provided in Sections 1.5 and 1.6. Motivated by Eloyan and Ghosh (2011), we propose a data reduction method in Section 1.5. We then present the PME algorithm in Section 1.6. A simulation study comparing the performance of the PME algorithm to some existing manifold learning methods is presented in Section 1.7. In Section 1.8, we propose a method for identifying the interiors of circle-like curves and cylinder/ball-like surfaces. The performance of the proposed method for estimating lung cancer tumor surfaces and interiors using computed tomography (CT) data is presented in Section 1.9.

1.2 Manifolds and Projection Indices

Before defining principal manifolds, we introduce concepts of manifolds and projection indices.

Definition 1.2.1. Let $f \in C(\mathbb{R}^d \rightarrow \mathbb{R}^D)$, then the image of f , i.e., $M_f^d = \{f(t) : t \in \mathbb{R}^d\}$, is called a d -dimensional manifold determined by f , where f is called an embedding map and $\mathbb{R}^d$ is called the parameter space. Furthermore, f is called a homeomorphism if its inverse $f^{-1} : M_f^d \rightarrow \mathbb{R}^d$ exists and is continuous. Here, the continuity of f^{-1} is associated with the subspace topology of M_f^d , i.e., the topology $\{U \cap M_f^d : U \text{ is an open subset of } \mathbb{R}^D\}$.

In applications, $\mathbb{R}^d$ is the space containing latent parameterizations $\{t_i\}_{i=1}^I$ of observed data $\{x_i\}_{i=1}^I \subset \mathbb{R}^D$.

The projection index π_f in (1.1.1) is well-defined under HS conditions when $d = 1$. We generalize π_f for all intrinsic dimensions d under a less stringent condition. Intuitively, the projection index of x to M_f^d is a parameter t such that $f(t)$ is closest to x (left panel of Figure 1.1). However, there might be more than one t such that $\|x - f(t)\|_{\mathbb{R}^D} = \inf_{t' \in \mathbb{R}^d} \|x - f(t')\|_{\mathbb{R}^D} =: \text{dist}(x, f)$, resulting in ambiguity in choosing t as shown in the right panel of Figure 1.1. To remove this ambiguity, we introduce the following function space.

$$C_\infty(\mathbb{R}^d \rightarrow \mathbb{R}^D) = \left\{ f \in C(\mathbb{R}^d \rightarrow \mathbb{R}^D) : \lim_{\|t\|_{\mathbb{R}^d} \rightarrow \infty} \|f(t)\|_{\mathbb{R}^D} = \infty \right\},$$

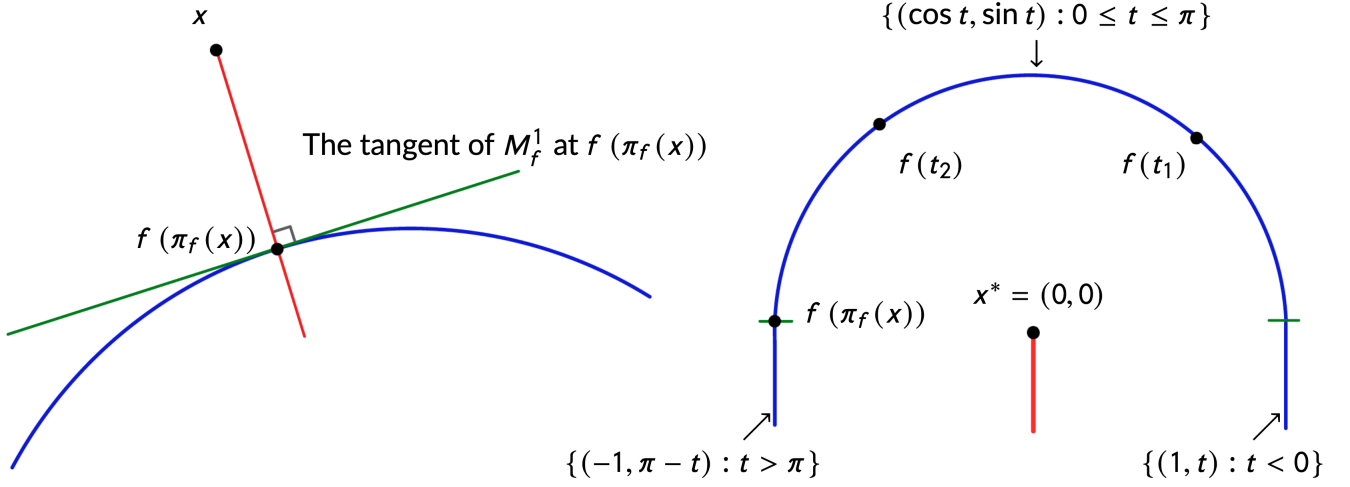


Figure 1.1: The left panel illustrates projection indices for $d = 1$. In the right panel, x^* is at the center of a semicircle. $\mathcal{A}_f(x^*) = [0, \pi]$ is compact, $\pi_f(x^*) = \pi$ and $\|x^* - f(t_1)\|_{\mathbb{R}^D} = \|x^* - f(t_2)\|_{\mathbb{R}^D} = \|x^* - f(\pi_f(x^*))\|_{\mathbb{R}^D} = \text{dist}(x^*, f)$ with $t_1 \neq t_2$. All the points in $\{(0, y) : y \leq 0\}$ (the red line) are ambiguity points.

In applications, this function space is not restrictive. Since a data set with finite sample size is always bounded, we are concerned with fitting functions in that bounded domain. Therefore the behavior of a C_∞ map as $\|t\|_{\mathbb{R}^d} \rightarrow \infty$ does not limit the applications of this framework. This approach is similar to focusing on the segment of a simple linear regression line within the range of an observed independent variable, even though the fitted line is unbounded. Based on these notations, we have the following theorem.

Theorem 1.2.1. *If $f \in C_\infty(\mathbb{R}^d \rightarrow \mathbb{R}^D)$, then the d -dimensional set*

$$\mathcal{A}_f(x) = \{t \in \mathbb{R}^d : \|x - f(t)\|_{\mathbb{R}^D} = \text{dist}(x, f)\}$$

is nonempty and compact for all $x \in \mathbb{R}^D$.

The proof of Theorem 1.2.1 is in the Appendix. The condition $f \in C_\infty(\mathbb{R}^d \rightarrow \mathbb{R}^D)$ is necessary for Theorem 1.2.1 to hold as shown by the following example. Let $f(t) = (\frac{e^t}{1+e^t}, 0)^T \notin C_\infty(\mathbb{R}^1 \rightarrow \mathbb{R}^2)$ and $x = (-1, 0)^T$, then $\inf_{t \in \mathbb{R}} \|f(t) - x\|_{\mathbb{R}^2} = \inf_{t \in \mathbb{R}} |\frac{e^t}{1+e^t} + 1| = 1$. However, $|\frac{e^t}{1+e^t} + 1| > 1$ for all t , then we have $\mathcal{A}_f(x) = \emptyset$. We propose a generalized definition of π_f as follows. Theorem 1.2.1 guarantees that generalized π_f is well-defined.

Definition 1.2.2. *Let $f \in C_\infty(\mathbb{R}^d \rightarrow \mathbb{R}^D)$. (i) If $\mathcal{A}_f(x)$ contains more than one element, then x is called an ambiguity point of f . (ii) For any $x \in \mathbb{R}^D$, the projection index $\pi_f(x)$ is defined as $\pi_f(x) = (t_1^*, t_2^*, \dots, t_d^*)$, where*

$$\begin{aligned} t_1^* &= \max \{t_1 : (t_1, t_2, \dots, t_d) \in \mathcal{A}_f(x)\}, \\ t_j^* &= \max \{t_j : (t_1^*, \dots, t_{j-1}^*, t_j, \dots, t_d) \in \mathcal{A}_f(x)\}, \quad j = 2, 3, \dots, d. \end{aligned} \quad (1.2.1)$$

The compactness of $\mathcal{A}_f(x)$ implies $\pi_f(x) \in \mathcal{A}_f(x)$. The projection index in (1.2.1) is a generalization of the

projection index defined in (1.1.1). Therefore, we use the same notation π_f to denote both projection indices. In defining the projection index (1.2.1), we replace the restrictive HS conditions with the less stringent condition $f \in C_\infty(\mathbb{R}^d \rightarrow \mathbb{R}^D)$ and allow d to be any positive integer. When $d = 1$, if $f \in C_\infty$ and f satisfies HS conditions, the projection index in (1.2.1) is the same as that in (1.1.1). The following theorem, whose proof is in Appendix, implies that $\pi_f(X)$ is a random d -vector if X is a random D -vector.

Theorem 1.2.2. *If $f \in C_\infty(\mathbb{R}^d \rightarrow \mathbb{R}^D)$, then (i) π_f is measurable, and (ii) $\{\pi_f(x) : x \in B\}$ is bounded when $B \subset \mathbb{R}^D$ is compact. Furthermore, if f has no ambiguity point in an open set $U \subset \mathbb{R}^D$ and is a homeomorphism, then (iii) π_f is continuous on U .*

1.3 Principal Manifolds

For a random D -vector X , its first d linear principal components are equivalent to the d -dimensional hyperplane defined by $\arg \min_{L \in \mathcal{L}} \mathbb{E} \|X - \Pi_L(X)\|_{\mathbb{R}^D}^2$, where $\mathcal{L}$ is the collection of all d -dimensional hyperplanes in $\mathbb{R}^D$ and $\Pi_L(X)$ is the projection of X to the hyperplane L . We propose a principal manifold framework generalizing PCA, replacing $\mathcal{L}$ with a Sobolev space. In the sequel, all derivatives are *generalized derivatives* defined on $\mathcal{D}'(\mathbb{R}^d)$, where $\mathcal{D}'(\mathbb{R}^d)$ is the collection of generalized functions on $\mathbb{R}^d$. Definitions of $\mathcal{D}'(\mathbb{R}^d)$ and generalized derivatives are provided in Chapter 6 of Rudin (1991). The only exception is that the derivatives referred to in the definition of the function space $C^k(\mathbb{R}^d \rightarrow \mathbb{R}^D)$ are still understood in the classical sense. Generalized derivatives, also called “derivatives of distributions,” apply to all locally integrable functions that may not be classically differentiable. They free us from smoothness assumptions in theoretical arguments.

1.3.1 Sobolev Spaces

We introduce the following function spaces of Sobolev type.

$$\begin{aligned} \nabla^{-\otimes 2} L^2(\mathbb{R}^d) &= \{f \in \mathcal{D}'(\mathbb{R}^d) : \|\nabla^{\otimes 2} f\|_{\mathbb{R}^{d \times d}} \in L^2(\mathbb{R}^d)\}, \\ H^2(\mathbb{R}^d) &= \{f \in L^2(\mathbb{R}^d) : \|\nabla f\|_{\mathbb{R}^d}, \|\nabla^{\otimes 2} f\|_{\mathbb{R}^{d \times d}} \in L^2(\mathbb{R}^d)\}, \\ \nabla^{-\otimes 2} L^2(\Omega) &= \{f|_\Omega : f \in \nabla^{-\otimes 2} L^2(\mathbb{R}^d)\}, \\ H^2(\Omega) &= \{f|_\Omega : f \in H^2(\mathbb{R}^d)\}, \end{aligned}$$

where Ω is any open subset of $\mathbb{R}^d$ with a smooth boundary, $f|_\Omega$ denotes the restriction of f to Ω , and

- ∇f denotes the vector-valued function $t \mapsto (\frac{\partial f}{\partial t_1}(t), \frac{\partial f}{\partial t_2}(t), \dots, \frac{\partial f}{\partial t_d}(t))^T =: \nabla f(t)$, i.e., the gradient of f , and $\|\nabla f\|_{\mathbb{R}^d}$ denotes the scalar-valued function $t \mapsto \|\nabla f(t)\|_{\mathbb{R}^d} = (\sum_{i=1}^d |\frac{\partial f}{\partial t_i}(t)|^2)^{1/2}$;
- $\nabla^{\otimes 2} f$ denotes the matrix-valued function $t \mapsto (\frac{\partial^2 f}{\partial t_i \partial t_j}(t))_{1 \leq i, j \leq d}$, i.e., the Hessian matrix of f ; $\|\nabla^{\otimes 2} f\|_{\mathbb{R}^{d \times d}}$

denotes the scalar-valued function $t \mapsto \|\nabla^{\otimes 2} f(t)\|_{\mathbb{R}^{d \times d}} = (\sum_{i,j=1}^d |\frac{\partial^2 f}{\partial t_i \partial t_j}(t)|^2)^{1/2}$.

- $\|\nabla f\|_{\mathbb{R}^d} \in L^2(\mathbb{R}^d)$ and $\|\nabla^{\otimes 2} f\|_{\mathbb{R}^d} \in L^2(\mathbb{R}^d)$ denote $\|\nabla f\|_{L^2(\mathbb{R}^d)}^2 = \int_{\mathbb{R}^d} \sum_{i=1}^d |\frac{\partial f}{\partial t_i}(t)|^2 dt < \infty$ and $\|\nabla^{\otimes 2} f\|_{L^2(\mathbb{R}^d)}^2 = \int_{\mathbb{R}^d} \sum_{i,j=1}^d |\frac{\partial^2 f}{\partial t_i \partial t_j}(t)|^2 dt < \infty$, respectively.
- $\nabla^{\otimes 2}$ is an abbreviation of the tensor product $\nabla \otimes \nabla = (\frac{\partial^2}{\partial t_i \partial t_j})_{i \leq i, j \leq d}$.

Section 1.5 of [Duchon \(1977\)](#) implies the following result.

Lemma 1.3.1. *If Ω is bounded, then $\nabla^{-\otimes 2} L^2(\Omega) = H^2(\Omega)$.*

If two functions are equal to each other almost everywhere with respect to the Lebesgue measure on $\mathbb{R}^d$, we identify them as the same function. Then we have the following regularity theorem.

Theorem 1.3.1. $\nabla^{-\otimes 2} L^2(\mathbb{R}^d) \subset C^k(\mathbb{R}^d)$, for $k < 2 - \frac{d}{2}$.

Proof. Let $f \in \nabla^{-\otimes 2} L^2(\mathbb{R}^d)$ and Ω be any open ball in $\mathbb{R}^d$, Lemma 1.3.1 implies $f|_{\Omega} \in H^2(\Omega)$. A Sobolev embedding theorem (Theorem 7.25 of [Rudin \(1991\)](#)) implies $f|_{\Omega} \in C^k(\Omega)$ for $k < 2 - \frac{d}{2}$. Then the result follows as Ω is arbitrary. $\square$

For simplicity, define $\nabla^{-\otimes 2} L^2(\mathbb{R}^d \rightarrow \mathbb{R}^D) = \{f(t) = (f_1(t), f_2(t), \dots, f_D(t))^T : f_l \in \nabla^{-\otimes 2} L^2(\mathbb{R}^d) \text{ for all } l = 1, 2, \dots, D\}$, and $C_{\infty} \cap \nabla^{-\otimes 2} L^2 = C_{\infty} \cap \nabla^{-\otimes 2} L^2(\mathbb{R}^d \rightarrow \mathbb{R}^D) = C_{\infty}(\mathbb{R}^d \rightarrow \mathbb{R}^D) \cap \nabla^{-\otimes 2} L^2(\mathbb{R}^d \rightarrow \mathbb{R}^D)$.

1.3.2 Definition of Principal Manifolds

As mentioned in Section 1.1, a problem of the model in [Kégl et al. \(2000\)](#), which is equivalent to (1.1.3) with $\|Pf\|_{\mathcal{H}} = \|f'\|_{L^2}$, is that the fitted f is not necessarily differentiable exactly everywhere. The main reason for this limitation is that the regularization from $\|f'\|_{L^2}^2$ may not provide enough penalty on the non-smoothness of f . We propose principal manifolds with higher regularity by replacing f' with the second derivative f'' . When $d = 1$ and f is arc-length parameterized, $\|f''(t)\|_{\mathbb{R}^D}$ is the curvature of f at t (Chapter 1.5 of [Do Carmo \(2016\)](#)). Additionally, $\int \|f''(t)\|_{\mathbb{R}^D}^2 dt$ is usually called the “total squared curvature” of f (e.g., [Enomoto and Okura \(2013\)](#)) and coincides with the penalty term of the cubic smoothing spline framework. For surfaces $f(t_1, t_2) = (f_1(t_1, t_2), \dots, f_D(t_1, t_2))^T$, the penalty $\sum_{l=1}^D \int_{\mathbb{R}^d} \sum_{i,j=1}^2 |\frac{\partial^2 f_l}{\partial t_i \partial t_j}(t)|^2 dt$ in the thin plate spline framework measures the roughness of surfaces M_f^2 (see [Koenker and Mizera \(2004\)](#)). For a general intrinsic dimension $d \geq 1$ and the map $f(t) = (f_1(t), f_2(t), \dots, f_D(t))^T$ with $t \in \mathbb{R}^d$, we generalize these two penalty terms to $\|\nabla^{\otimes 2} f\|_{L^2(\mathbb{R}^d)}^2 = \sum_{l=1}^D \|\nabla^{\otimes 2} f_l\|_{L^2(\mathbb{R}^d)}^2 = \sum_{l=1}^D \int_{\mathbb{R}^d} \sum_{i,j=1}^d |\frac{\partial^2 f_l}{\partial t_i \partial t_j}(t)|^2 dt$. In the sequel, we apply $\|\nabla^{\otimes 2} f\|_{L^2(\mathbb{R}^d)}^2$ to measure the roughness of M_f^d . We broaden the use of the name “total squared curvature” defined for curves and call its high-dimensional generalization $\|\nabla^{\otimes 2} f\|_{L^2(\mathbb{R}^d)}^2$ *total squared curvature* as well. The tolerance of large total squared curvature increases the complexity and decreases the stability of fitted manifolds. Therefore, we penalize fitted manifolds with large total squared curvature. These considerations motivate us to define principal manifolds as follows.

Definition 1.3.1. Let X be a random D -vector associated with the probability measure or density function $\mathbb{P}$ such that X has compact support $\text{supp}(\mathbb{P})$ and finite second moments. Let $f, g \in C_\infty \cap \nabla^{-\otimes 2} L^2(\mathbb{R}^d \rightarrow \mathbb{R}^D)$ and $\lambda \in [0, \infty]$, we define the following functionals

$$\mathcal{K}_{\lambda, \mathbb{P}}(f, g) = \mathbb{E} \|X - f(\pi_g(X))\|_{\mathbb{R}^D}^2 + \lambda \|\nabla^{\otimes 2} f\|_{L^2(\mathbb{R}^d)}^2, \quad \mathcal{K}_{\lambda, \mathbb{P}}(f) = \mathcal{K}_{\lambda, \mathbb{P}}(f, f). \quad (1.3.1)$$

A manifold $M_{f^*}^d$ determined by f^* is called a *principal manifold* for X (or $\mathbb{P}$) with the tuning parameter λ if

$$f_\lambda^* = \arg \min_{f \in \mathcal{F}(\mathbb{P})} \mathcal{K}_{\lambda, \mathbb{P}}(f), \quad (1.3.2)$$

$$\text{where } \mathcal{F}(\mathbb{P}) = \left\{ f \in C_\infty \cap \nabla^{-\otimes 2} L^2(\mathbb{R}^d \rightarrow \mathbb{R}^D) : \sup_{x \in \text{supp}(\mathbb{P})} \|\pi_f(x)\|_{\mathbb{R}^d} = 1 \right\}.$$

Since λ above is allowed to be ∞ , we adopt the convention $\infty \times 0 = 0$ as $\lim_{\lambda \rightarrow \infty} (\lambda \times 0) = 0$. Then $\mathcal{K}_{\infty, \mathbb{P}}(f) < \infty$ only if $\|\nabla^{\otimes 2} f\|_{L^2(\mathbb{R}^d)} = 0$. Suppose f_λ^* is derived, then $\pi_{f_\lambda^*}(X)$ gives a d -dimensional parameterization of X . Theorem 1.2.2 (iii) implies the continuity of $\pi_{f_\lambda^*}(X)$ under some conditions on f_λ^* . The constraint $\sup\{\|\pi_f(x)\|_{\mathbb{R}^d} : x \in \text{supp}(\mathbb{P})\} = 1$ is well-defined (see Theorem 1.2.2 (ii)) and restricts the parameterizations $\{\pi_f(x) : x \in \text{supp}(\mathbb{P})\}$ exactly in the unit ball $\{t \in \mathbb{R}^d : \|t\|_{\mathbb{R}^d} \leq 1\}$. This restriction is an analog of the arc-length parameterization in the scenario $d = 1$. Additionally, Theorem 1.3.1 implies the regularity of principal manifolds, i.e., $f_\lambda^* \in C^k(\mathbb{R}^d \rightarrow \mathbb{R}^D)$ for $k < \max\{2 - \frac{d}{2}, 1\}$. The definition above is a special case of the regularized principal manifolds (1.1.3) defined by Smola et al. (2001), where $P = \nabla^{\otimes 2}$, $\mathcal{H} = L^2(\mathbb{R}^d)$, and $\mathcal{F} = \mathcal{F}(\mathbb{P})$.

Motivated by the “projection-adaptation” algorithm in Section 5 of Smola et al. (2001), we apply its iterative fashion to estimate principal manifolds. Specifically, we estimate $f_\lambda^* = \arg \min_{f \in \mathcal{F}(\mathbb{P})} \mathcal{K}_{\lambda, \mathbb{P}}(f)$ using

$$f_{(n)} = \arg \min_f \{\mathcal{K}_{\lambda, \mathbb{P}}(f, f_{(n-1)}) : f \in C_\infty \cap \nabla^{-\otimes 2} L^2(\mathbb{R}^d \rightarrow \mathbb{R}^D)\}, \quad n = 1, 2, \dots, \quad \lambda \geq 0. \quad (1.3.3)$$

Suppose (1.3.3) stops when $n = n^*$, and we obtain $f_{(n^*)} \in C_\infty \cap \nabla^{-\otimes 2} L^2(\mathbb{R}^d \rightarrow \mathbb{R}^D)$. Then an estimate of $\arg \min_{f \in \mathcal{F}(\mathbb{P})} \mathcal{K}_{\lambda, \mathbb{P}}(f)$ is given by $\widehat{f_\lambda^*}(t) = f_{(n^*)}(\kappa t)$ with $\kappa = \sup\{\|\pi_{f_{(n^*)}}(x)\|_{\mathbb{R}^d} : x \in \text{supp}(\mathbb{P})\} < \infty$ (see Theorem 1.2.2 (ii)), and $\widehat{f_\lambda^*} \in \mathcal{F}(\mathbb{P})$. Computing $\pi_{f_{(n)}}(X)$ in $\mathcal{K}_{\lambda, \mathbb{P}}(f, f_{(n)})$ implicitly corresponds to the “projection” step discussed in Section 1.2, where our results guarantee that $\pi_{f_{(n)}}(X)$ is well-defined. Minimizing $\mathcal{K}_{\lambda, \mathbb{P}}(f, f_{(n)})$ with respect to f corresponds to the “adaptation” step. Since $f_{(n)} \in C_\infty \cap \nabla^{-\otimes 2} L^2$ for all n , Theorem 1.3.1 guarantees the regularity of $f_{(n)}$ for all n . The iteration (1.3.3) usually approximates local minima. Hence, successful implementation of the iteration depends on the choice of the starting values. An initialization strategy with a partial implementation of ISOMAP is illustrated in Section 1.7.

1.3.3 Two Special Cases

In this subsection, we discuss two extreme special cases of the tuning parameter λ : $\lambda = \infty$ and $\lambda = 0$. We show that these two cases imply linear PCA and the HS principal curve algorithm, respectively. Besides, we discuss potential issues that may arise when using these two extreme cases in applications, leading to the consideration of other values of λ in our proposed framework. The following theorem establishes the fact that $\lambda = \infty$ implies linear PCA.

Theorem 1.3.2. *Suppose X is a random D -vector with finite second moments, $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_D$ and $e_1, e_2, \dots, e_D$ are eigenvectors and eigenvalues of the covariance matrix of X , respectively. $\mathbf{v}_i$ corresponds to e_i and $e_1 \geq \dots \geq e_d > e_{d+1} \geq \dots \geq e_D$. Then the principal manifold for X with tuning parameter $\lambda = \infty$ is the linear manifold $\left\{ \mathbb{E}X + \sum_{i=1}^d \alpha_i \mathbf{v}_i : \alpha_i \in \mathbb{R}^1 \right\}$.*

Its proof is in the Appendix. Theorem 1.3.2 implies that a large λ shrinks principal manifolds towards PCA. However, if the underlying manifolds are nonlinear, the linear manifolds with zero curvature are not satisfactory estimators.

When $\lambda = 0$, the estimation of principal manifolds may result in overfitting. For example, let $\mathbb{P}$ be the empirical distribution $\frac{1}{I} \sum_{i=1}^I \delta_{x_i}$ for the data $\{x_i\}_{i=1}^I$. For any $f \in C^\infty(\mathbb{R}^d \rightarrow \mathbb{R}^D)$ satisfying $\{x_i\}_{i=1}^I \subset M_f^d$, i.e., f passes through every data point, $\sup_i \|\pi_f(x_i)\|_{\mathbb{R}^d} = 1$, and $f(t)$ is equal to an affine function when $\|t\|_{\mathbb{R}^d} > M$ for a sufficiently large $M > 0$, we have $\mathcal{K}_{0,\mathbb{P}}(f) = 0$. Then M_f^d is a principal manifold for $\mathbb{P}$ and overfits the data $\{x_i\}_{i=1}^I$. Additionally, $\lambda = 0$ results in self-consistency (1.1.2), potentially resulting in the saddle issue discussed in Section 1. Specifically, the HS principal curve algorithm is of the following form.

$$f_{(n)} = \mathcal{T}_{HS} f_{(n-1)}, \quad \text{where } \mathcal{T}_{HS} f_{(n-1)}(t) = \mathbb{E}(X | \pi_{f_{(n-1)}}(X) = t), \quad n = 1, 2, \dots \quad (1.3.4)$$

If $f_{(n)}$ converges to f , then (1.3.4) implies (1.1.2). The following theorem implies that (1.3.4) is a special case of (1.3.3) with $\lambda = 0$.

Theorem 1.3.3. *If $f_{(n-1)}$ and $\mathcal{T}_{HS} f_{(n-1)} \in C_\infty \cap \nabla^{-\otimes 2} L^2$, then*

$$\mathcal{T}_{HS} f_{(n-1)} = \arg \min_f \left\{ \mathcal{K}_{0,\mathbb{P}}(f, f_{(n-1)}) : f \in C_\infty \cap \nabla^{-\otimes 2} L^2 \right\}.$$

Proof. Let $\mathcal{M}$ be the collection of measurable $\mathbb{R}^d \rightarrow \mathbb{R}^D$ maps. We have $\inf \{ \mathcal{K}_{0,\mathbb{P}}(f, f_{(n-1)}) : f \in C_\infty \cap \nabla^{-\otimes 2} L^2 \} \geq \inf \{ \mathbb{E} \|X - f(\pi_{f_{(n-1)}}(X))\|_{\mathbb{R}^D}^2 : f \in \mathcal{M} \} = \mathbb{E} \|X - \mathbb{E}(X | \pi_{f_{(n-1)}}(X))\|_{\mathbb{R}^D}^2 = \mathbb{E} \|X - \mathcal{T}_{HS} f_{(n)}(\pi_{f_{(n)}}(X))\|_{\mathbb{R}^D}^2$. Then $\mathcal{T}_{HS} f_{(n-1)} \in C_\infty \cap \nabla^{-\otimes 2} L^2(\mathbb{R}^d \rightarrow \mathbb{R}^D)$ implies the desired result. $\square$

It is important to consider the case when λ is positive but very close to 0. In this case, the total squared curvature $\|\nabla^{\otimes 2} f_\lambda^*\|_{L^2(\mathbb{R}^d)}^2$ of $f_\lambda^* = \arg \min_{f \in \mathcal{F}(\mathbb{P})} \mathcal{K}_{\lambda,\mathbb{P}}(f)$ may be unreasonably large. The large value of the total

squared curvature may result in numerical issues when performing optimization procedures. When $\lambda > 0$, similar to the penalty of the ridge regression, $\|\nabla^{\otimes 2} f_{\lambda}^*\|_{L^2(\mathbb{R}^d)}^2$ is bounded by a positive U_{λ} , a number that depends on λ . As λ decreases, the bound U_{λ} tends to increase. Therefore, in some cases, we may have $\limsup_{\lambda \rightarrow 0} \|\nabla^{\otimes 2} f_{\lambda}^*\|_{L^2(\mathbb{R}^d)}^2 = \infty$. We call this phenomenon “curvature singularity at $\lambda = 0$ ” throughout this paper. If the curvature singularity at $\lambda = 0$ holds, a fitted manifold $M_{f_{\lambda}^*}^d$ may have high roughness and be unstable for sufficiently small values of λ . Results in this subsection imply that $\lambda = \infty$ may mask the underlying curvature of the function, $\lambda = 0$ may result in overfitting or the saddle issue, and small values of λ may result in the curvature singularity at $\lambda = 0$.

1.4 A New Approach to Principal Manifold Estimation

We present the outline of our proposed novel principal manifold estimation (PME) algorithm, briefly describing its three main steps - reduction, fitting, and tuning. This algorithm addresses two problems in the framework described in Section 1.3: (i) the choice of $\lambda \in (0, \infty)$, and (ii) the reduction of the computational burden in implementing iteration (1.3.3) and elimination of effects of outliers. In image analysis applications, the size of data is usually very large, resulting in computational burden when applying manifold learning algorithms. One approach to addressing computational burden is subsampling, e.g., Yue et al. (2016). While leading to faster computation, subsampling may result in removing important sections of a given data set. Our proposed PME algorithm solves these two problems simultaneously.

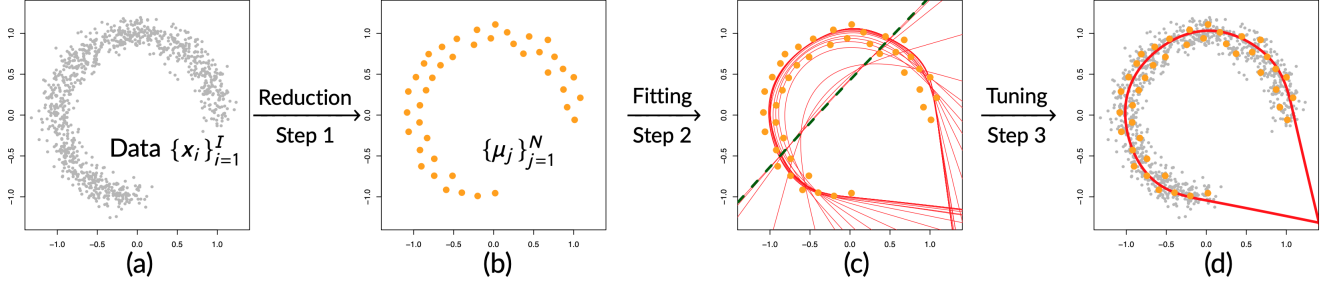


Figure 1.2: (a) Data $\{x_i\}_{i=1}^I$ (gray). (b) Each dot (orange) denotes a μ_j . (c) Estimated $\hat{f}_{\lambda}$ (red curves) are associated with different $\lambda > 0$. The straight line (green and dashed) is associated with $\lambda \rightarrow \infty$ and is an approximation of the first principal component. (d) Using $\{x_i\}_{i=1}^I$, we choose the optimal λ^* and draw the corresponding $\hat{f}_{\lambda^*}$ (red curve).

A step-by-step visualization of the proposed algorithm is in Figure 1.2. In the first stage, the sample size I of the data $\{x_i\}_{i=1}^I$ from $\mathbb{P}$ is reduced to obtain a collection of points $\{\mu_j\}_{j=1}^N$ with a smaller size N , where each μ_j is associated with a weight θ_j satisfying $\sum_{j=1}^N \theta_j = 1$, such that $\{\mu_j\}_{j=1}^N$ preserve the geometric features of the underlying manifold and are less noisy than the original sample $\{x_i\}_{i=1}^I$. Then $\mathcal{K}_{\lambda, \hat{Q}_N}(f) = \sum_{j=1}^N \theta_j \|\mu_j - f(\pi_f(\mu_j))\|_{\mathbb{R}^D}^2 + \lambda \|\nabla^{\otimes 2} f\|_{L^2(\mathbb{R}^d)}^2$ approximates $\mathcal{K}_{\lambda, \mathbb{P}}(f) \approx \frac{1}{I} \sum_{i=1}^I \|x_i - f(\pi_f(x_i))\|_{\mathbb{R}^D}^2 + \lambda \|\nabla^{\otimes 2} f\|_{L^2(\mathbb{R}^d)}^2$ (see Theorem 1.5.3 in Section 1.5.3), where $\mathcal{K}_{\lambda, \hat{Q}_N}(f)$ is the functional in (1.3.1) associated with the probability measure $\hat{Q}_N =$

$\sum_{j=1}^N \theta_j \delta_{\mu_j}$. This stage results in the reduction of computational burden and elimination of effects of outliers (see Figure 1.3 (c) and (d)). In the second stage of this approach, for a preselected set of tuning parameters $\lambda > 0$, we estimate $\hat{f}_\lambda = \arg \min_f \mathcal{K}_{\lambda, \hat{Q}_N}(f)$. The collection of estimated functions $\{\hat{f}_\lambda\}_{\lambda > 0}$ prevents the curvature singularity at $\lambda = 0$. Finally, in the third stage, we choose an optimal tuning parameter λ^* that preserves the geometric structure in the data while avoiding overfitting $\{\mu_j\}_{j=1}^N$. The following theorem implies that $\{\hat{f}_\lambda\}_{\lambda > 0}$ prevents the curvature singularity at $\lambda = 0$, i.e, $\limsup_{\lambda \rightarrow 0} \|\nabla^{\otimes 2} \hat{f}_\lambda\|_{L^2(\mathbb{R}^d)}^2 < \infty$.

Theorem 1.4.1. *For all $\lambda > 0$, let $\hat{f}_\lambda = \arg \min_f \{\mathcal{K}_{\lambda, \hat{Q}_N}(f) : f \in \mathcal{F}(\hat{Q}_N)\}$ with $\hat{Q}_N = \sum_{j=1}^N \theta_j \delta_{\mu_j}$. Then we have the bound*

$$\begin{aligned} & \sup \left\{ \|\nabla^{\otimes 2} \hat{f}_\lambda\|_{L^2(\mathbb{R}^d)}^2 : \lambda > 0 \right\} \\ & \leq \inf \left\{ \|\nabla^{\otimes 2} f\|_{L^2(\mathbb{R}^d)}^2 : f \in \mathcal{F}(\hat{Q}_N), \text{ and } f(\pi_f(\mu_j)) = \mu_j \text{ for } j = 1, 2, \dots, N \right\} =: U_N < \infty. \end{aligned}$$

Proof. $\mathcal{W}_N = \{f \in \mathcal{F}(\hat{Q}_N) : f(\pi_f(\mu_j)) = \mu_j \text{ for } j = 1, 2, \dots, N\} \neq \emptyset$ implies $U_N < \infty$. If $\sup_{\lambda > 0} \|\nabla^{\otimes 2} \hat{f}_\lambda\|_{L^2(\mathbb{R}^d)}^2 > U_N$, there exist $\tilde{\lambda} > 0$ and $\tilde{f} \in \mathcal{W}_N$ such that $\|\nabla^{\otimes 2} \hat{f}_{\tilde{\lambda}}\|_{L^2(\mathbb{R}^d)}^2 > \|\nabla^{\otimes 2} \tilde{f}\|_{L^2(\mathbb{R}^d)}^2$. Then $\mathcal{K}_{\tilde{\lambda}, \hat{Q}_N}(\tilde{f}) = \tilde{\lambda} \|\nabla^{\otimes 2} \tilde{f}\|_{L^2(\mathbb{R}^d)}^2 < \mathcal{K}_{\tilde{\lambda}, \hat{Q}_N}(\hat{f}_{\tilde{\lambda}})$, which contradicts the definition of $\hat{f}_{\tilde{\lambda}} = \arg \min_{f \in \mathcal{F}(\hat{Q}_N)} \mathcal{K}_{\tilde{\lambda}, \hat{Q}_N}(f)$. $\square$

1.5 Step 1: Data Reduction

The reduction step of the PME algorithm is motivated by the data generating mechanism in manifold learning tasks. In this Section, we show that the distribution of the data $\{x_i\}_{i=1}^I$ can be approximated by a probability measure of the form $\hat{Q}_N = \sum_{j=1}^N \theta_j \delta_{\mu_j}$, a convex combination of the point masses δ_{μ_j} . We propose the *high dimensional mixture density estimation* (HDMDE) algorithm for the estimation of μ_j , θ_j , and N in $\hat{Q}_N$.

1.5.1 Motivation and Estimation of Mixture Density Parameters

In manifold learning, we assume that the D -dimensional data $\{x_i\}_{i=1}^I$ are realizations from a d -dimensional latent manifold, corrupted by D -dimensional noise. Each x_i is generated in two stages - the latent data stage and the noise corruption stage. In the latent data stage, a latent random D -vector T is generated from a probability measure Q^* , where Q^* is supported on a d -dimensional manifold. Then in the noise corruption stage, given $T = t$, the data point x_i is generated from a probability density function (PDF) $\psi(\cdot - t)$ with $\int x \psi(x) dx = \mathbf{0} \in \mathbb{R}^D$. Then the distribution generating x_i is the PDF $p(x) = \psi * Q^*(x) = \int \psi(x - t) Q^*(dt)$, where $*$ denotes the convolution operation. We may estimate the latent probability measure Q^* by maximizing the nonparametric likelihood $\mathcal{L}(Q) = \prod_{i=1}^I \psi * Q(x_i)$ in $Q \in \mathcal{Q}$, where $\mathcal{Q}$ denotes the collection of probability measures supported on d -dimensional manifolds. Theorem 3.1 of Lindsay (1983) implies that there exists a unique probability measure of the form $\hat{Q}_N = \sum_{j=1}^N \theta_j \delta_{\mu_j}$ with $N \leq I$ achieving $\sup_{Q \in \mathcal{Q}} \mathcal{L}(Q)$. For example, in Figure 1.2 (d), the gray dots denote data x_i , the red curve denotes

the support of the latent variable T (the probability measure Q^*), and the large orange dots illustrate the point masses δ_{μ_j} in $\hat{Q}_N$.

Substituting the true Q^* with the maximizer $\hat{Q}_N$, we estimate the distribution generating x_i by the PDF $\psi * \hat{Q}_N(x) = \sum_{j=1}^N \theta_j \psi(x - \mu_j)$. To estimate μ_j , θ_j , and N , we use a mixture density estimation approach. For any $\sigma > 0$, denote $\psi_\sigma(x) = \frac{1}{\sigma^D} \psi\left(\frac{x}{\sigma}\right)$. For any fixed positive integer N , let $\{\mu_{j,N}\}_{j=1}^N$ be a collection of points in $\mathbb{R}^D$ depending on N , let $\theta_N = (\theta_{1,N}, \theta_{2,N}, \dots, \theta_{N,N})^T$ be in the probability simplex $\Theta_N = \{\theta_N : \theta_{j,N} \geq 0, \sum_{j=1}^N \theta_{j,N} = 1\}$, and let σ_N be a positive number such that $\lim_{N \rightarrow \infty} \sigma_N = 0$. We construct the following mixture density

$$p_N(x|\theta_N) = \psi_{\sigma_N} * \hat{Q}_N(x) = \sum_{j=1}^N \theta_{j,N} \psi_{\sigma_N}(x - \mu_{j,N}), \quad \text{where } \hat{Q}_N = \sum_{j=1}^N \theta_{j,N} \delta_{\mu_{j,N}}. \quad (1.5.1)$$

We estimate $\mu_{j,N}$, θ_N , σ_N , and N such that $p_N(x|\theta_N)$ approximates the true PDF $p(x)$ generating data $\{x_i\}_{i=1}^I$.

Assuming that the number of mixture components N is fixed, various approaches may be implemented for the estimation of mixture parameters $\mu_{j,N}$, $\theta_{j,N}$, and σ_N . A common approach for estimating these parameters is based on the EM algorithm (Dempster et al. (1977)). However, this approach can be too computationally intensive in our setting. We propose a high-dimensional generalization of the mixture density estimation algorithm proposed by Eloyan and Ghosh (2011), where the estimation of $\mu_{j,N}$ and σ_N is performed in a computationally efficient manner for a given N and the estimation of the mixture weights $\theta_{j,N}$ is then conducted using the EM algorithm.

Estimation of $\mu_{j,N}$: Partition $\{x_i\}_{i=1}^I$ to N clusters by *k-means clustering*. Let $\{\mu_{j,N}\}_{j=1}^N$ be the centers of the N clusters.

Estimation of σ_N : Let $\{x_{j,l}\}_{l=1}^{L_j}$ be x_i in j^{th} cluster. We estimate σ_N by

$$\hat{\sigma}_N = \left\{ \frac{1}{D \times N} \sum_{j=1}^N \left(\frac{1}{L_j} \sum_{l=1}^{L_j} \|x_{j,l} - \mu_{j,N}\|_{\mathbb{R}^D}^2 \right) \right\}^{1/2}.$$

If $\{x_{j,l}\}_{l=1}^{L_j}$ are iid $N_D(\mu_{j,N}, \sigma_N^2 I_{D \times D})$ for $j = 1, 2, \dots, N$, then $\hat{\sigma}_N^2$ is an unbiased estimator of σ_N^2 .

Estimation of $\theta_{j,N}$: Assuming that $\{x_i\}_{i=1}^I$ is a random sample from the PDF

$$p_N(x|\theta_N) = \sum_{j=1}^N \theta_{j,N} \psi_{\sigma_N}(x - \mu_{j,N}) \approx p(x),$$

we estimate $\theta_{j,N}$ by likelihood maximization. In practice, the sample mean is used as an unbiased estimate of $\mathbb{E}_{p_N}(X|\theta_N)$. Therefore, we apply the constraint $\int_{\mathbb{R}^D} x p_N(x|\theta_N) dx = \bar{x} = \frac{1}{I} \sum_{i=1}^I x_i$ and estimate $\theta_{j,N}$ achieving the likelihood maximization using a constrained EM algorithm. The detailed derivation of this procedure is in

Appendix. The proposed approach results in the following equation for updating the estimate of $\theta_{j,N}$.

$$\begin{aligned}\theta_{j,N}^{(k)} &= \frac{\sum_{i=1}^I w_{ij} \left(\theta_N^{(k-1)} \right)}{\hat{\rho}_1 + \hat{\rho}_2^T \mu_{j,N}}, \quad j = 1, 2, \dots, N \text{ and } k = 1, 2, \dots, \text{ where} \\ (\hat{\rho}_1, \hat{\rho}_2) &= \arg \min_{\rho_1 \in \mathbb{R}, \rho_2 \in \mathbb{R}^D} \left\{ \left\| \sum_{j=1}^N \left(\frac{\sum_{i=1}^I w_{ij} \left(\theta_N^{(k-1)} \right)}{\rho_1 + \rho_2^T \mu_{j,N}} \right) - 1 \right\|^2 + \left\| \sum_{j=1}^N \left(\frac{\sum_{i=1}^I w_{ij} \left(\theta_N^{(k-1)} \right)}{\rho_1 + \rho_2^T \mu_{j,N}} \right) \mu_{j,N} - \bar{x} \right\|_{\mathbb{R}^D}^2 \right\}, \\ w_{ij}(\theta_N) &= \frac{\theta_{j,N} \times \psi_{\hat{\sigma}_N}(x_i - \mu_{j,N})}{\sum_{j'=1}^N \theta_{j',N} \times \psi_{\hat{\sigma}_N}(x_i - \mu_{j',N})}.\end{aligned}\tag{1.5.2}$$

If $\sup \left\{ |\theta_{j,N}^{(k^*)} - \theta_{j,N}^{(k^*-1)}| : j = 1, 2, \dots, N \right\}$ is smaller than a predetermined threshold, then the iteration (1.5.2) stops, and $\theta_N^{(k^*)} = (\theta_{1,N}^{(k^*)}, \theta_{2,N}^{(k^*)}, \dots, \theta_{N,N}^{(k^*)})^T$ is an estimate of θ_N .

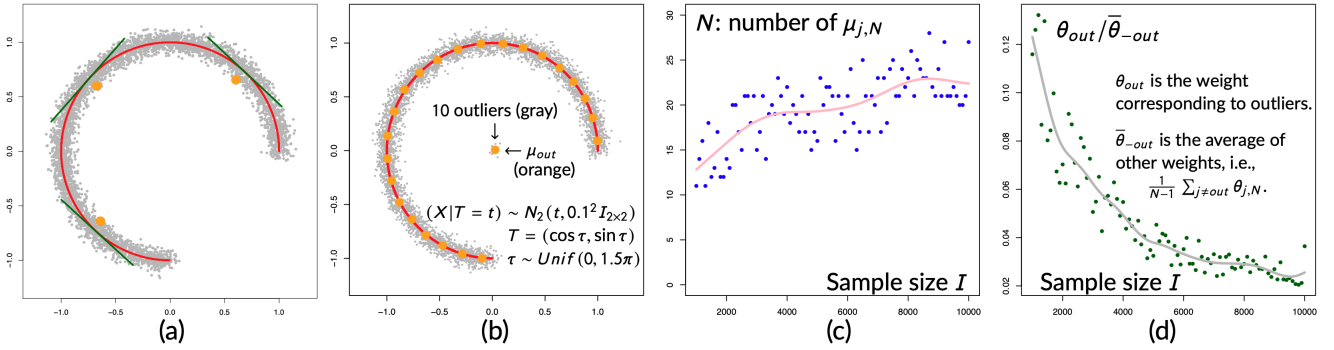


Figure 1.3: (a) The red 3/4 circle is the underlying manifold of interest, the data cloud (gray) is generated from the X in (b), applying k -means clustering with $N = 3$ results in three large centers (orange); the 3 centers lie completely inside the circle and do not give a good representation of the red circle. This issue stems from the nonzero curvature: For each tangent (green) of the circle, data are unevenly distributed on its two sides. Cluster centers are “dragged” to the side with more data points. This issue disappears when N is sufficiently large so that $\|p_N(\cdot|\theta_N) - p\|_{L^1(\mathbb{R}^D)} \approx 0$. (b) Small dots (gray) are random samples from X . The support of T is the solid curve (red). We apply Algorithm 1 to these gray dots with input $N_0 = 10, \alpha = 0.05, \epsilon = 0.001$. The large dots (orange) denote the estimated $\mu_{j,N}$ in the $\hat{Q}_N = \sum_{j=1}^N \theta_{j,N} \delta_{\mu_{j,N}}$. (c) Set $N_0 = 10$, for random samples with size I ranging from 1000 to 10000, the estimated N are shown by dots (blue). The curve (pink) shows the trend of N as the sample size I increases. (d) Illustration of the influence of outliers on $\hat{Q}_N$ as I increases. For each I , the influence of outliers is measured by the quantity $\theta_{out}/\bar{\theta}_{-out}$ and shown by a dot (green). The gray curve shows the corresponding trend.

1.5.2 Estimation of the Number of Mixture Components

In this subsection, we propose an iterative hypothesis testing procedure to choose the number of mixture components N . If N is too small, $\hat{Q}_N$ in (1.5.1) may not capture geometric features of data x_i . A detailed example of the “small N ” issue is presented in Figure 1.3 (a). On the other hand, an unreasonably large N may result in computational burden and redundant model complexity. Motivated by the following theorem, we choose N by investigating the L^1 -distance between $p_N(\cdot|\theta_N) = \psi * \hat{Q}_N$ in (1.5.1) and the PDF p generating data x_i . The diameter of a set U in $\mathbb{R}^D$ is defined by $\text{diam}(U) = \sup\{\|x_1 - x_2\|_{\mathbb{R}^D} : x_1, x_2 \in U\}$.

Theorem 1.5.1. Suppose p is a PDF with compact support $\text{supp}(p) = \overline{\{x : p(x) \neq 0\}}$ and $p \in L^q(\mathbb{R}^D)$ for some $1 \leq q < \infty$, $\mathcal{M}_N = \{\mu_{j,N}\}_{j=1}^N \subset \text{supp}(p)$, and $d_N, \sigma_N > 0$ for all positive integers N . If (i) the triplet (d_N, σ_N, ψ) satisfies

$$\lim_{N \rightarrow \infty} \left(\sup \left\{ \|\psi_{\sigma_N}(\cdot - y) - \psi_{\sigma_N}\|_{L^q(\mathbb{R}^D)} : \|y\|_{\mathbb{R}^D} \leq d_N \right\} \right) = \lim_{N \rightarrow \infty} \sigma_N = \lim_{N \rightarrow \infty} d_N = 0; \quad (1.5.3)$$

(ii) there exists a partition of the compact set $\text{supp}(p)$, i.e., $\text{supp}(p) = \bigcup_{j=1}^N A_{j,N}$ with $A_{i,N} \cap A_{j,N} = \emptyset$ when $i \neq j$, such that $A_{j,N} \cap \mathcal{M}_N = \{\mu_{j,N}\}$ and $\sup\{\text{diam}(A_{j,N}) : j = 1, 2, \dots, N\} \leq d_N$ for all large positive integers N ; then there exists a sequence $\{\theta_N\}_N$ with $\theta_N \in \Theta_N$ such that $\lim_{N \rightarrow \infty} \|p_N(\cdot|\theta_N) - p\|_{L^q(\mathbb{R}^D)} = 0$, where $p_N(\cdot|\theta_N)$ is defined by (1.5.1).

The proof of Theorem 1.5.1 is in Appendix. In applications, observed data are always in a bounded domain. Thus the assumption on p is not restrictive. The triplet satisfying (1.5.3) exists, e.g., ψ is any PDF, $\sigma_N = N^{-\alpha_1}$, and $d_N = N^{-(\alpha_1 + \alpha_2)}$, then this triplet satisfies (1.5.3) when $q = 1$, where α_1 and α_2 are allowed to be any positive numbers. Herein, we set ψ to be the standard Gaussian kernel $(2\pi)^{-D/2} \exp\{-\|x\|_{\mathbb{R}^D}^2/2\}$. Condition (ii) essentially requires $\mathcal{M}_N$ to be dense in $\text{supp}(p)$ as $N \rightarrow \infty$. Since we estimate the knots $\mu_{j,N}$ as centers of the N k-means clusters of the data $\{x_i\}_{i=1}^I \sim_{iid} p$, $\mathcal{M}_N = \{\mu_{j,N}\}_{j=1}^N$ tends to be dense in $\text{supp}(p)$ as the number of clusters increases. Therefore, condition (ii) is realistic. Since all PDFs are in $L^1(\mathbb{R}^D)$, we are only interested in the special case of Theorem 1.5.1 where $q = 1$.

The limit $\lim_{N \rightarrow \infty} \|p_N(\cdot|\theta_N) - p\|_{L^1(\mathbb{R}^D)} = 0$ implies $\lim_{N \rightarrow \infty} \|p_{N+1}(\cdot|\theta_{N+1}) - p_N(\cdot|\theta_N)\|_{L^1(\mathbb{R}^D)} = 0$. If we further assume $p \in L^\infty(\mathbb{R}^D)$, then we have the following limit motivating the proposed method of selecting N .

$$|\mathbb{E}_p \{p_{N+1}(X|\theta_{N+1}) - p_N(X|\theta_N)\}| \leq \|p\|_{L^\infty(\mathbb{R}^D)} \|p_{N+1}(\cdot|\theta_{N+1}) - p_N(\cdot|\theta_N)\|_{L^1(\mathbb{R}^D)} \rightarrow 0, \quad \text{as } N \rightarrow \infty,$$

where $X \sim p$. This limit implies $\mathbb{E}_p \{p_{N+1}(X|\theta_{N+1}) - p_N(X|\theta_N)\} \approx 0$ when N is sufficiently large. Therefore, we choose a sufficiently large N by testing the following hypothesis.

$$H_0 : \mathbb{E}_p \{p_{N+1}(X|\theta_{N+1}) - p_N(X|\theta_N)\} = 0 \quad \text{vs} \quad H_a : \mathbb{E}_p \{p_{N+1}(X|\theta_{N+1}) - p_N(X|\theta_N)\} \neq 0, \quad (1.5.4)$$

where $\mathbb{E}_p$ is the expectation associated with the PDF p . Since p and θ_N are unknown, we use $\bar{\Delta}_{I,N} = \frac{1}{I} \sum_{i=1}^I \hat{\Delta}_i$ to test the hypothesis (1.5.4), where $\hat{\theta}_N = \hat{\theta}_N(X_1, X_2, \dots, X_I)$ is an estimator of θ_N computed from the iid sample X_i and $\hat{\Delta}_i = p_{N+1}(X_i|\hat{\theta}_{N+1}) - p_N(X_i|\hat{\theta}_N)$. We apply the following theorem to perform an asymptotic test of (1.5.4).

Theorem 1.5.2. Suppose $\hat{\theta}_n$ is an estimator of the true $\theta_n \in \Theta_n$, such that $\hat{\theta}_n = \theta_n + o_p(I^{-1/2})$, where $n \in \{N, N+1\}$, N is fixed, and $\psi \in L^\infty(\mathbb{R}^D)$. Denote $\Delta_i = p_{N+1}(X_i|\theta_{N+1}) - p_N(X_i|\theta_N)$, $\mu_{\Delta,N} = \mathbb{E}_p \Delta_1$, $\hat{S}_{I,N}^2 = \frac{1}{I} \sum_{i=1}^I \hat{\Delta}_i^2 - (\bar{\Delta}_{I,N})^2$ and $\hat{S}_{I,N} = \sqrt{\hat{S}_{I,N}^2}$. Then $\sqrt{I} \frac{\bar{\Delta}_{I,N} - \mu_{\Delta,N}}{\hat{S}_{I,N}} \rightarrow N(0, 1)$ in distribution as $I \rightarrow \infty$.

Theorem 1.5.2 can be derived directly from the central limit theorem and Slutsky's theorem, hence its proof is omitted. Since we are interested in testing the hypothesis $H_0 : \mu_{\Delta,N} = 0$ as shown in (1.5.4), we define the statistic $Z_{I,N} = \sqrt{I} \frac{\bar{\Delta}_{I,N}}{\bar{S}_{I,N}}$. From Theorem 1.5.2, under H_0 , we have $Z_{I,N} \sim N(0,1)$ approximately when I is large. We choose N by

$$N = N_\alpha = \inf \{N : |Z_{I,N}| < z_{1-\alpha/2} \text{ and } N_0 \leq N \leq I-1\},$$

where N_0 denotes a predetermined lower bound for N , $z_{1-\alpha/2}$ is the $1 - \alpha/2$ quantile of $N(0,1)$, and $\alpha = 0.05$ is chosen for testing (1.5.4). Figure 1.3 (c) shows that the estimated N tends to be much smaller than I . In both the method proposed above and the counterpart in Eloyan and Ghosh (2011), the number of mixture components N is chosen by measuring the dissimilarity between $p_{N+1}(\cdot|\theta_{N+1})$ and $p_N(\cdot|\theta_N)$. Eloyan and Ghosh (2011) applies *Kullback-Leibler divergence* (KLD) while we apply L^1 -norm. We choose L^1 -norm because we found in simulations that L^1 -norm captures the geometric features of p better than KLD.

1.5.3 The High-Dimensional Mixture Density Estimation (HDMDE) Algorithm

By iteratively combining the procedures for the estimation of mixture weights $\theta_{j,N}$, mixture element means $\mu_{j,N}$, and the number of mixture components N presented in Sections 1.5.1 and 1.5.2, we propose the HDMDE algorithm in Algorithm 1 to estimate $\hat{Q}_N = \sum_{j=1}^N \theta_{j,N} \delta_{\mu_{j,N}}$. In this Section, we use simulation studies to illustrate that HDMDE results in a substantial reduction in computation speed and elimination of the effects of outliers. In addition, we present a proof-of-concept simulation study showing that the density estimated by HDMDE can approximate the true density better than the kernel density estimate (KDE) in terms of minimizing the L^1 -distance.

Algorithm 1 HDMDE

Input: (i) Data points $\{x_i\}_{i=1}^I$ in $\mathbb{R}^D$, (ii) a positive integer $N_0 < I-1$, and (iii) $\epsilon, \alpha \in (0,1)$.

Output: N , $\{\mu_{j,N}\}_{j=1}^N$, $\hat{\theta}_N = (\hat{\theta}_{1,N}, \dots, \hat{\theta}_{N,N})^T$, and σ_N . Then we have

$$\hat{Q}_N = \sum_{j=1}^N \hat{\theta}_{j,N} \delta_{\mu_{j,N}} \text{ and } p_N(\cdot|\hat{\theta}_N) = \psi_{\sigma_N} * \hat{Q}_N.$$

- 1: $N \leftarrow N_0$ and formally $Z_{I,N} \leftarrow 2 \times z_{1-\alpha/2}$.
- 2: Estimate $\mu_{j,N}$ and σ_N using the k-means clustering.
- 3: Apply the iteration (1.5.2) and get a sequence $\{\hat{\theta}_N^{(k)}\}_k$. Set $\hat{\theta}_N = \hat{\theta}_N^{(k^*)}$ with

$$k^* = \min \left\{ k : \sup_j \left| \hat{\theta}_{j,N}^{(k-1)} - \hat{\theta}_{j,N}^{(k)} \right| < \epsilon \right\}.$$

- 4: Compute $p_N(x_i|\hat{\theta}_N) = \sum_{j=1}^N \hat{\theta}_{j,N} \times \psi_{\sigma_N}(x_i - \mu_{j,N})$ for all i .
 - 5: **while** $|Z_{I,N}| \geq z_{1-\alpha/2}$ and $N < I-1$, **do**
 - 6: $N \leftarrow N+1$, repeat the steps 2, 3, 4, and compute $Z_{I,N}$.
 - 7: **end while**
-

When the sample size I of data $\{x_i\}_{i=1}^I$ is large, manifold fitting can be computationally expensive. One advan-

tage of using HDMDE in PME is the comparatively small computational burden in estimating $\hat{f}_\lambda = \arg \min_f \mathcal{K}_{\lambda, \hat{Q}_N}(f)$. We conduct simulation studies to compare the magnitudes of estimated N and sample size I empirically and show that N is much smaller than I . Figure 1.3 (b,c) provides an illustrative comparison between N and I by simulations. For each I ranging from 1000 to 10000, we generate $I - 10$ points close to a $3/4$ part of a unit circle as presented in Figure 1.3 (b) and 10 outliers from $N_2(\mathbf{0}, 0.1^2 I_{2 \times 2})$. We estimate a $\hat{Q}_N$ by HDMDE for each simulated sample. In Figure 1.3 (b), we show one simulation example with $I = 5000$ by gray points and the estimated $\{\mu_{j,N}\}_{j=1}^N$ by large orange dots. Figure 1.3 (c) illustrates the estimated N versus I and shows that N are much smaller than I .

Another advantage of HDMDE is that the effect of outliers on $\hat{Q}_N$ is negligible. As a result, when we fit a manifold by $\hat{f}_\lambda = \arg \min_f \mathcal{K}_{\lambda, \hat{Q}_N}(f)$, the result is robust to outliers. Specifically, if the node $\mu_{j',N}$ is closer to outliers than to the main part of the data cloud, the associated weight $\theta_{j',N}$ is small. In each of the simulations in Figure 1.3, only one node defined as μ_{out} is located in the outlier cluster, i.e., in $\{x : \|x\|_{\mathbb{R}^2} < 0.3\}$. We denote the weight associated with μ_{out} by θ_{out} , and denote the average of other weights $\frac{1}{N-1} \sum_{j \neq out} \theta_{j,N}$ by $\bar{\theta}_{-out}$. The ratio $\theta_{out}/\bar{\theta}_{-out}$ measures the influence of μ_{out} compared to that of other $\mu_{j,N}$. The lower this ratio, the more negligible the effect of outliers on estimation of $\hat{Q}_N$. Figure 1.3 (d) shows that $\theta_{out}/\bar{\theta}_{-out}$ is small and decreases drastically as the sample size I increases. Hence, the point μ_{out} representing 10 outliers has a negligible effect on $\hat{Q}_N$ and this effect decreases as I increases.

An important property of HDMDE is its performance in approximating the true PDF p in terms of minimizing the L^1 -distance. We compare HDMDE with KDE in terms of approximating p and then use this property of HDMDE to justify the reduction from $\mathcal{K}_{\lambda,p}(f)$ to $\mathcal{K}_{\lambda, \hat{Q}_N}(f)$ in PME. To apply KDE in a simulation example, we implement the R function `kde` in the package `ks` using default parameter values provided in the package. Let $p_{hdmde}(\cdot|\mathbf{X})$ and $p_{kde}(\cdot|\mathbf{X})$ denote the PDFs estimated by applying HDMDE and KDE, respectively, to data $\mathbf{X} = \{X_i\}_{i=1}^I \sim_{iid} p$. The difference between the performances of HDMDE and KDE is measured by $\|p_{kde}(\cdot|\mathbf{X}) - p\|_{L^1(\mathbb{R}^D)} - \|p_{hdmde}(\cdot|\mathbf{X}) - p\|_{L^1(\mathbb{R}^D)} =: \mathcal{J}(\mathbf{X})$. Let p be the PDF of the random vector X in Figure 1.3 (b) (without the 10 outliers). Using this PDF p as an example, we generate 500 realizations of $\mathbf{X}$ from p with $I = 1000$ and estimate the mean $\mathbb{E}\mathcal{J}(\mathbf{X})$ and variance $\mathbb{V}\mathcal{J}(\mathbf{X})$ using the sample mean and sample variance. Then we compute the Wald 95%-confidence interval $\mathbb{E}\mathcal{J}(\mathbf{X}) \pm 1.96\sqrt{\mathbb{V}\mathcal{J}(\mathbf{X})} \approx (0.018, 0.130)$. This interval shows that, on average, HDMDE performs better than KDE in the L^1 -approximation of the density p in this example. Using this property of HDMDE, we provide the following result showing that minimizing $\mathcal{K}_{\lambda,p}(f)$ is approximately equivalent to minimizing $\mathcal{K}_{\lambda, \hat{Q}_N}(f)$.

Theorem 1.5.3. *Suppose (i) p and ψ are PDFs with bounded support; (ii) $\{\hat{Q}_N = \sum_{j=1}^N \theta_{j,N} \delta_{\mu_{j,N}}\}_{N=1}^\infty$ satisfies $\{\mu_{j,N}\}_{j=1}^N \subset \text{supp}(p)$ for all N , and $\lim_{N \rightarrow \infty} \|\psi_{\sigma_N} * \hat{Q}_N - p\|_{L^1(\mathbb{R}^d)} = 0$ for a sequence $\{\sigma_N\}_{N=1}^\infty$ with $\lim_{N \rightarrow \infty} \sigma_N = 0$; (iii) $f \in C_\infty \cap \nabla^{-\otimes 2} L^2$ is a homeomorphism and has no ambiguity point in a neighborhood of $\text{supp}(p)$. If there exists $\{\mu_j\}_{j=1}^\infty \subset \mathbb{R}^D$ so that $\lim_{N \rightarrow \infty} \mu_{j,N} = \mu_j$ and $\sum_{j=1}^\infty (\sup_{N': N' \geq j} \theta_{j,N'}) < \infty$, we have the limit $\lim_{N \rightarrow \infty} \mathcal{K}_{\lambda, \hat{Q}_N}(f) = \mathcal{K}_{\lambda,p}(f)$ for $\lambda \in [0, \infty]$.*

The proof of Theorem 1.5.3 is in Appendix. Although the Gaussian kernel ψ does not have a bounded support, most of its mass is in a bounded domain, e.g., the Gaussian kernel in $\mathbb{R}^3$ satisfies $\psi(x) \leq 10^{-22}$ when $\|x\|_{\mathbb{R}^3} \geq 10$. In Theorem 1.5.3, condition (ii) can be implied by Theorem 1.5.1, and condition (iii) is related to Theorem 1.2.2.

1.6 Step 2: Fitting, Step 3: Tuning, Model Complexity Selection

In this Section, we propose the details of Step 2 (fitting) and Step 3 (tuning) of our proposed PME algorithm illustrated in Figure 1.2. To fit $\hat{f}_\lambda = \arg \min_f \mathcal{K}_{\lambda, \hat{Q}_N}(f)$ in Step 2 of the PME algorithm, we apply the iteration (1.3.3) with $\mathbb{P} = \hat{Q}_N$, i.e., $f_{(n+1)} = \arg \min_f \left\{ \mathcal{K}_{\lambda, \hat{Q}_N}(f, f_{(n)}) : f = (f_1, f_2, \dots, f_D)^T \in C_\infty \cap \nabla^{-\otimes 2} L^2(\mathbb{R}^d \rightarrow \mathbb{R}^D) \right\}$ with

$$\mathcal{K}_{\lambda, \hat{Q}_N}(f, f_{(n)}) = \sum_{l=1}^D \left\{ \sum_{j=1}^N \theta_{j,N} \left| \mu_{j,N,l} - f_l(\pi_{f_{(n)}}(\mu_{j,N})) \right|^2 + \lambda \left\| \nabla^{\otimes 2} f_l \right\|_{L^2(\mathbb{R}^d)}^2 \right\}, \quad (1.6.1)$$

where $\mu_{j,N,l}$ is the l^{th} component of the D -vector $\mu_{j,N}$, and l denotes a vector component index. Define the following notations: (i) If ν is an even integer, $\eta_\nu(t) = \|t\|_{\mathbb{R}^d}^\nu \log(\|t\|_{\mathbb{R}^d})$ when $\|t\|_{\mathbb{R}^d} \neq 0$ and $\eta_\nu(t) = 0$ when $\|t\|_{\mathbb{R}^d} = 0$; otherwise, $\eta_\nu(t) = \|t\|_{\mathbb{R}^d}^\nu$. (ii) $Poly_1[t]$ is the linear space of polynomials on $\mathbb{R}^d$ with degree ≤ 1 and has a linear basis $\{p_k\}_{k=1}^{d+1}$. The following theorem implies that the minimizer of $\mathcal{K}_{\lambda, \hat{Q}_N}(\cdot, f_{(n)})$ in $C_\infty \cap \nabla^{-\otimes 2} L^2$ is of a spline form.

Theorem 1.6.1. *Suppose $f_{(n)} \in C_\infty(\mathbb{R}^d \rightarrow \mathbb{R}^D)$, $d \leq 3$, and each polynomial in $Poly_1[t]$ is uniquely determined by its values on $\mathcal{C} = \{\pi_{f_{(n)}}(\mu_{j,N})\}_{j=1}^N$. Then a minimizer of $\mathcal{K}_{\lambda, \hat{Q}_N}(\cdot, f_{(n)})$ within $C_\infty \cap \nabla^{-\otimes 2} L^2(\mathbb{R}^d \rightarrow \mathbb{R}^D)$ is of the following form.*

$$f_{(n+1),l}(t) = \sum_{j=1}^N s_{j,l} \times \eta_{4-d}(t - \pi_{f_{(n)}}(\mu_{j,N})) + \sum_{k=1}^{d+1} \alpha_{k,l} \times p_k(t), \quad l = 1, 2, \dots, D, \quad (1.6.2)$$

with the constraint $\sum_{j=1}^N s_{j,l} \times p_k(\pi_{f_{(n)}}(\mu_{j,N})) = 0$ for all $k = 1, 2, \dots, d+1$ and $l = 1, 2, \dots, D$.

Proof of Theorem 1.6.1 is in Appendix. The reason for the dimension restriction $d \leq 3$ is that $\nabla^{-\otimes 2} L^2(\mathbb{R}^d)$ is a reproducing kernel Hilbert space only if $d \leq 3$ (see Wahba (1990), Chapter 2.4). For the purpose of visualization, the intrinsic dimension $d \leq 3$ is not restrictive. When $d = 1$, (1.6.2) is a cubic smoothing spline. When $d = 2$, (1.6.2) is a thin plate spline. From Theorem 1.6.1 and the calculation strategy in Chapter 2 of Wahba (1990), it follows that minimizing (1.6.1) in $C_\infty \cap \nabla^{-\otimes 2} L^2$ is equivalent to obtaining the minimizers

$$\arg \min_{(s_l, \alpha_l)} \left\{ \left\| \mathbf{W}^{1/2} (\mu_l - \mathbf{E} s_l - \mathbf{T} \alpha_l) \right\|_{\mathbb{R}^N}^2 + \lambda \left\| \mathbf{E}^{1/2} s_l \right\|_{\mathbb{R}^N}^2 : s_l \in \mathbb{R}^N, \alpha_l \in \mathbb{R}^{d+1}, \text{ and } \mathbf{T}^T s_l = 0 \right\}$$

for $l = 1, 2, \dots, D$, where

- $\mathbf{T}$ is an $N \times (d+1)$ matrix whose $(i, j)^{th}$ element is $p_j(\pi_{f(n)}(\mu_{i,N}))$;
- $\mu_l = (\mu_{1,N,l}, \mu_{2,N,l}, \dots, \mu_{N,N,l})^T$, $\alpha_l = (\alpha_{1,l}, \alpha_{2,l}, \dots, \alpha_{d+1,l})^T$, $s_l = (s_{1,l}, s_{2,l}, \dots, s_{N,l})^T$ for $l = 1, 2, \dots, D$;
- $\mathbf{E}$ is an $N \times N$ matrix whose $(i, j)^{th}$ element is $\eta_{4-d}(\pi_{f(n)}(\mu_{i,N}) - \pi_{f(n)}(\mu_{j,N}))$;
- $\mathbf{W} = \text{diag}(\theta_{1,N}, \theta_{2,N}, \dots, \theta_{N,N})$.

Using the Lagrange multiplier method, we can obtain these minimizers by solving the following linear equations.

$$\begin{pmatrix} 2\mathbf{E}\mathbf{W}\mathbf{E} + 2\lambda\mathbf{E} & 2\mathbf{E}\mathbf{W}\mathbf{T} & \mathbf{T} \\ 2\mathbf{T}^T\mathbf{W}\mathbf{E} & 2\mathbf{T}^T\mathbf{W}\mathbf{T} & \mathbf{0} \\ \mathbf{T}^T & \mathbf{0} & \mathbf{0} \end{pmatrix} \begin{pmatrix} s_l \\ \alpha_l \\ m_l \end{pmatrix} = \begin{pmatrix} 2\mathbf{E}\mathbf{W}\mu_l \\ 2\mathbf{T}^T\mathbf{W}\mu_l \\ \mathbf{0} \end{pmatrix}, \quad l = 1, 2, \dots, D, \quad (1.6.3)$$

where m_l are Lagrange multipliers. The coefficient matrix in (1.6.3) is symmetric, has many zero elements, and of order $N+2d+2$. Since N is moderate in most applications (see Figure 1.3 (c)), solving (1.6.3) is not computationally expensive.

As detailed in (1.6.1), we use $\hat{Q}_N$ to estimate $\hat{f}_\lambda$ for each $\lambda > 0$. This procedure shrinks the collection of candidate functions from $C_\infty \cap \nabla^{-\otimes 2} L^2$ to the one-parameter family $\{\hat{f}_\lambda\}_{\lambda>0}$. Theorem 1.4.1 shows that this approach prevents curvature singularity at $\lambda = 0$. Then we choose an optimal element $\hat{f}_{\lambda^*}$ in $\{\hat{f}_\lambda\}_{\lambda>0}$ by using the observed data $\{x_i\}_{i=1}^I$ to tune the family $\{\hat{f}_\lambda\}_{\lambda>0}$. Specifically, we choose $\hat{f}_{\lambda^*}$ which minimizes the MSD $\mathcal{D}(\hat{f}_\lambda)$ associated with $\{x_i\}_{i=1}^I$, i.e.,

$$\lambda^* = \arg \min_{\lambda>0} \left\{ \mathcal{D}(\hat{f}_\lambda) \right\}, \quad \text{where } \mathcal{D}(\hat{f}_\lambda) = \frac{1}{I} \sum_{i=1}^I \left\| x_i - \hat{f}_\lambda \left(\pi_{\hat{f}_\lambda}(x_i) \right) \right\|_{\mathbb{R}^D}^2. \quad (1.6.4)$$

In applications, higher values of the tuning parameter λ reduce the effect of corrupting noise. The reduction from $\{x_i\}_{i=1}^I$ to $\hat{Q}_N$ reduces the corrupting noise and, hence, the corresponding λ is expected to be small. Therefore, the estimated optimal λ^* tends to be small. Figure 1.4 illustrates the relationships between $\log \lambda$, $\log \lambda^*$, and MSD $\mathcal{D}(\hat{f}_\lambda)$.

Determining the pair (N, λ^*) by HDMDE and (1.6.4) completes our model complexity selection procedure and, hence, the PME algorithm. The PME algorithm presented in Algorithm 2 encapsulates the procedures presented in Sections 1.5 and 1.6. The R code for performing estimation using Algorithm 2 is available at <https://github.com/KMengBrown/Principal-Manifold-Estimation.git>. While a rigorous proof of the convergence of Algorithm 2 is outside of the scope of this paper, the algorithm converged in almost all of the simulation studies conducted.

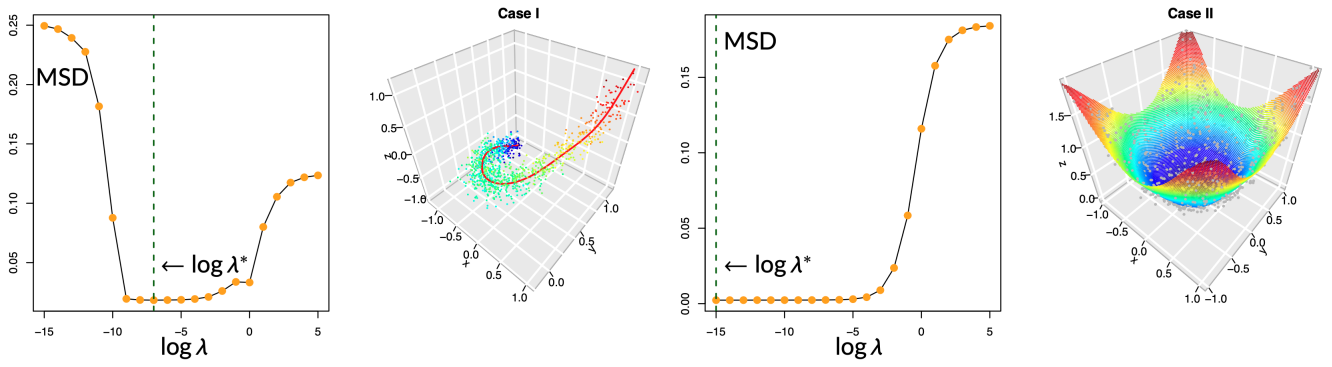


Figure 1.4: Points (colored) in Case I are 1000 realizations of X with $(X|T = t) \sim N_3(t, 0.1I_{3 \times 3})$, $\tau \sim Unif(-1, 1)$, $T = (\tau, \tau^2, \tau^3)^T$. The points (gray) in Case II are 1000 realizations of X with $(X|T = t) \sim N_3(t, 0.05I_{3 \times 3})$, $\tau_1, \tau_2 \sim iid Unif(-1, 1)$, $T = (\tau_1, \tau_2, \tau_1^2 + \tau_2^2)^T$. Using Algorithm 2, we fit the points in Case I and Case II, respectively. With tuning parameters $\lambda = e^k, k = -15, -14, \dots, 5$, we plot MSD versus $\log \lambda$ as above. The (green) dash lines indicate optimal tuning parameters. For Case II, the reason why the smallest λ is chosen is that the corresponding reduced points $\mu_{j,N}$ (with associated weights $\theta_{j,N}$) have almost no information of the 3-dimensional corrupting noise $N(t, 0.05I_{3 \times 3})$.

1.7 Simulations

We compare the PME algorithm to existing methods for simulated data in the following three scenarios with dimension pairs $(d = 1, D = 2)$, $(d = 1, D = 3)$, and $(d = 2, D = 3)$. Simulation analyses in this section are implemented in the R software (R Core Team (2019)). In the implementation of PME (Algorithm 2), we set its inputs as follows: $\lambda_g = \exp(g)$ for $g = -15, -14, \dots, 5$, $N_0 = 20 \times D$, $\alpha = 0.05$, $\epsilon = 0.001$, $\epsilon^* = 0$, *itr* values are given in Tables 1.1 and 1.2. For the first two dimension pairs, we compare PME to two methods: (i) The HS principal curve algorithm using the R function `principal_curve` in package `princurve` (version 2.1.4). Three `smoother` options - `smooth.spline`, `lowess`, and `periodic_lowess` - are provided in this R function. In each simulation, we apply all the three smoothers and apply the one producing the smallest MSD $\mathcal{D}(f) = \frac{1}{I} \sum_{i=1}^I \|x_i - f(\pi_f(x_i))\|_{\mathbb{R}^D}^2$. (ii) ISOMAP: Parameterize the I data points $\{x_i\}_{i=1}^I$, by $\{t_i\}_{i=1}^I$ using ISOMAP and then fit a map $f_{isomap}^* = \arg \min_f \{ \frac{1}{I} \sum_{i=1}^I \|x_i - f(t_i)\|_{\mathbb{R}^D}^2 + \lambda^* \|\nabla^{\otimes 2} f\|_{L^2(\mathbb{R}^1)}^2 : f \in \nabla^{-\otimes 2} L^2(\mathbb{R}^1 \rightarrow \mathbb{R}^D) \}$, where λ^* is chosen to be the optimal tuning parameter obtained in the corresponding PME fit to make the comparison fair. The minimizer f_{isomap}^* is achieved by cubic smoothing splines (Duchon (1977)). For $(d = 2, D = 3)$, we compare PME to two methods: (i) the principal surface (PS) algorithm introduced by Yue et al. (2016), where the optimal number of basis functions in PS is obtained by the new cross-validation method proposed by Yue et al. (2016), using the R function for PS provided by the first author of Yue et al. (2016); (ii) ISOMAP: Conducted as described above for its counterpart with $d = 1$, except we apply thin plate splines to achieve the corresponding minimizer (Duchon (1977)). For all scenarios, the performance measurement of a fitted f is the MSD $\mathcal{D}(f)$. Although we apply ISOMAP to the parameterization step of PME, the implementation of PME has an important advantage compared to solely applying ISOMAP. When directly applying ISOMAP to the observed data, we use the full set of I data points for

Algorithm 2 PME Algorithm:

Input: (i) Data points $\{x_i\}_{i=1}^I$ in $\mathbb{R}^D$; (ii) a positive integer $N_0 < I - 1$; (iii) $\alpha, \epsilon, \epsilon^* \in (0, 1)$; (iv) candidate tuning parameters $\{\lambda_g\}_{g=1}^G$; (v) $itr \geq 1$, which is the maximum number of iterations allowed.

Output: (i) Analytic formula of $f^* : \mathbb{R}^d \rightarrow \mathbb{R}^D$ determining the fitted manifold $M_{f^*}^d$; (ii) optimal tuning parameter λ^* .

- 1: Apply HDMDE (Algorithm 1) with input $(\{x_i\}_{i=1}^I, N_0, \epsilon, \alpha)$ and obtain N , $\{\mu_{j,N}\}_{j=1}^N$, and $\{\theta_{j,N}\}_{j=1}^N$.
- 2: **Parameterization:** Apply ISOMAP to parameterize the reduced collection $\{\mu_{j,N}\}_{j=1}^N$ by the d -dimensional parameters $\{t_j\}_{j=1}^N$. Formally set $\pi_{f_{(0)}}(\mu_{j,N}) \leftarrow t_j$ for $j = 1, 2, \dots, N$.
- 3: **for all** $g = 1, 2, \dots, G$ **do**
- 4: $\lambda \leftarrow \lambda_g$ and obtain $f_{(1)}$ by solving (1.6.3).
- 5: $\mathcal{E} \leftarrow 2 \times \epsilon^*$, $n \leftarrow 1$, and $\mathcal{D}_{\hat{Q}_N}(f_{(1)}) \leftarrow \sum_{j=1}^N \theta_{j,N} \|\mu_{j,N} - f_{(1)}(\pi_{f_{(1)}}(\mu_{j,N}))\|_{\mathbb{R}^D}^2$.
- 6: **while** $\mathcal{E} \geq \epsilon^*$ and $n < itr$, **do**
- 7: Compute $f_{(n+1)}$ from $f_{(n)}$ by solving (1.6.3) and

$$\mathcal{D}_{\hat{Q}_N}(f_{(n+1)}) \leftarrow \sum_{j=1}^N \theta_{j,N} \|\mu_{j,N} - f_{(n+1)}(\pi_{f_{(n+1)}}(\mu_{j,N}))\|_{\mathbb{R}^D}^2.$$

- 8: $\mathcal{E} \leftarrow |\mathcal{D}_{\hat{Q}_N}(f_{(n+1)}) - \mathcal{D}_{\hat{Q}_N}(f_{(n)})| / \mathcal{D}_{\hat{Q}_N}(f_{(n)})$ and $n \leftarrow n + 1$.
- 9: **end while**
- 10: $\hat{f}_g \leftarrow f_{(n)}$.
- 11: **end for**
- 12: $\kappa \leftarrow \max\{\|\pi_{\hat{f}_g}(x_i)\|_{\mathbb{R}^d} : i = 1, 2, \dots, I\}$, where

$$g^* = \arg \min_g \left\{ \frac{1}{I} \sum_{i=1}^I \left\| x_i - \hat{f}_g \left(\pi_{\hat{f}_g}(x_i) \right) \right\|_{\mathbb{R}^D}^2 : g = 1, 2, \dots, G \right\}.$$

- 13: $f^*(t) = \hat{f}_{g^*}(\kappa t)$, where the analytic formula of f^* is from (1.6.2), and $\lambda^* \leftarrow \lambda_{g^*}$.
-

parameterization, while within PME we apply ISOMAP to parameterize the reduced set $\{\mu_{j,N}\}_{j=1}^N$ in the process of iteration. Since ISOMAP is computationally expensive for high-dimensional data, and the sample size I tends to be much larger than N (see Figure 1.3 (c)), directly applying ISOMAP to the full observed data is much more time consuming than applying the PME approach. As shown by our simulation studies, this gain in reduction of computation time is accompanied by similar performance of PME as compared to ISOMAP.

Visualizations of the estimation results for some curves and surfaces are presented in Figures 1.5 and 1.6. When $(d = 1, D = 2)$, we generate data under four settings in Figure 1.5 and apply all corresponding methods for each of the four cases; when $(d = 1, D = 3)$ or $(d = 2, D = 3)$, we generate data using the three different data cloud scenarios in Figure 1.6 and apply all corresponding methods in each of all the three cases. For each method in each case, we run 100 simulations with simulated data sets of size $I = 1000$ and summarize the simulation results in Tables 1.1 and 1.2, where the mean and standard deviation (sd) of the 100 simulated MSD in each case for each method are presented. The column “itr” in these tables shows the number of iterations conducted for each algorithm. Column groups (a) (b) (c) (d) in Table 1.1 correspond to the panels (a) (b) (c) (d) in Figure 1.5, respectively. Column groups (a) (b) in Table 1.2 correspond to the panels (a) (b) in Figure 1.6, respectively. The column group (c) in Table 1.2 corresponds to the row (c), i.e., the lower panels, of Figure 1.6. Except for ISOMAP, all methods take less than ten minutes to run in each simulation in all cases on a PC with a 2.6 GHz Intel Core

i5 processor and 8 GB 1600 MHz DDR3 memory. In all our simulations, when $d = 1$, PME and HS take a similar amount of time to run; when $d = 2$, PME and PS take a similar amount of time to run. Further optimization of authors' R code should make the proposed PME algorithm more efficient.

Table 1.1: MSD comparison: $d = 1$ and $D = 2$. (The unit of mean and sd is 10^{-3} , and the lowest mean in each column is in bold.)

Methods	(a)			(b)			(c)			(d)		
	itr	mean	sd	itr	mean	sd	itr	mean	sd	itr	mean	sd
PME	20	5.995	0.4082	100	40.77	1.682	10	10.09	0.6029	5	23.66	1.082
HS	300	28.33	8.690	200	351.8	8.702	100	12.96	0.4344	5	24.21	1.120
ISOMAP	0	5.712	0.3297	0	40.97	1.802	0	10.12	0.4171	0	23.50	0.8739

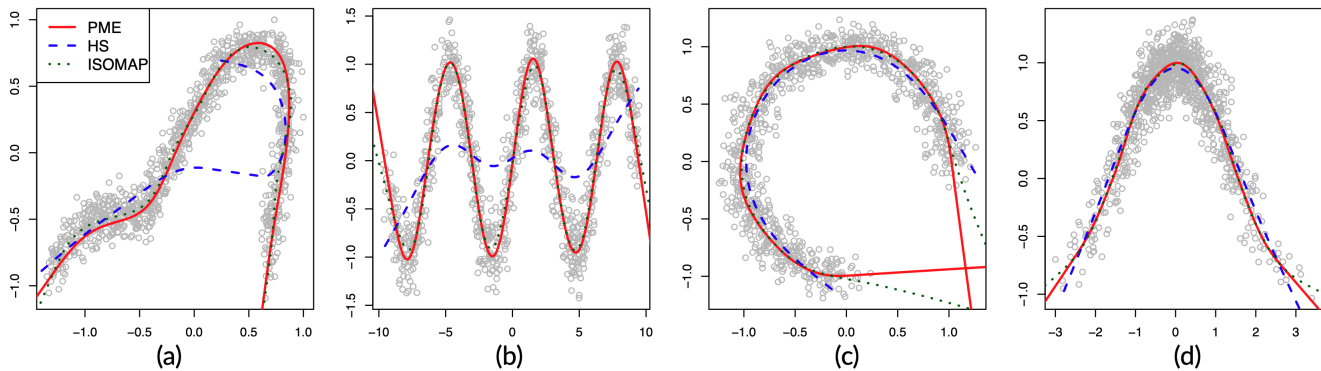


Figure 1.5: Illustration of simulation settings. In each setting, data (in gray) are generated as follows: (a) a 1/4 part of one slice of a CT data set presented in Section 1.9 is used with added Gaussian noise; (b) realizations of X with $(X|T=t) \sim N_2(t, 0.2I_{2 \times 2})$, $T = (\tau, \sin \tau)^T$ and $\tau \sim \text{Unif}(-3\pi, 3\pi)$; (c) realizations of the X in Figure 1.3 (b) (without the 10 outliers); (d) realizations of X with $(X|T=t) \sim N_2(t, 0.15I_{2 \times 2})$, $T = (\tau, \cos \tau)^T$ and $\tau \sim N(0, 1)$.

Table 1.2: MSD comparison: $d = 1, 2$ and $D = 3$. (The unit of mean and sd is 10^{-3} , and the lowest mean in each column is in bold.)

$d = 1$		(a)			(b)			$d = 2$		(c)		
Methods		itr	mean	sd	itr	mean	sd	Methods	itr	mean	sd	
PME		100	18.58	0.6023	100	5.320	0.2115	PME	10	2.522	0.1138	
HS		200	21.23	0.6294	500	88.03	0.7479	PS	10	2.520	0.1137	
ISOMAP		0	19.52	0.6163	0	5.214	0.1661	ISOMAP	0	2.496	0.1103	

Simulation results: (i) For $(d = 1, D = 2)$, Figure 1.5 (a) (b) show that PME performs much better than HS. Figure 1.5 (c) (d) show that PME performs slightly better than HS. ISOMAP and PME perform similarly well in all four cases. The difference between them is visible only near the tails of data clouds. (ii) For $(d = 1, D = 3)$, Figure 1.6 (a) shows that the three methods perform similarly well. Figure 1.6 (b) shows that PME and ISOMAP perform similarly well, and both of them perform much better than HS. (iii) For $(d = 2, D = 3)$, Figure 1.6 (c) shows that PME, PS, and ISOMAP perform equally well. Tables 1.1 and 1.2 support these conclusions. In conclusion, PME performs either significantly or marginally better than HS, and PME is not inferior to ISOMAP. However, ISOMAP is extremely time consuming in all scenarios compared to other methods. If we increase the

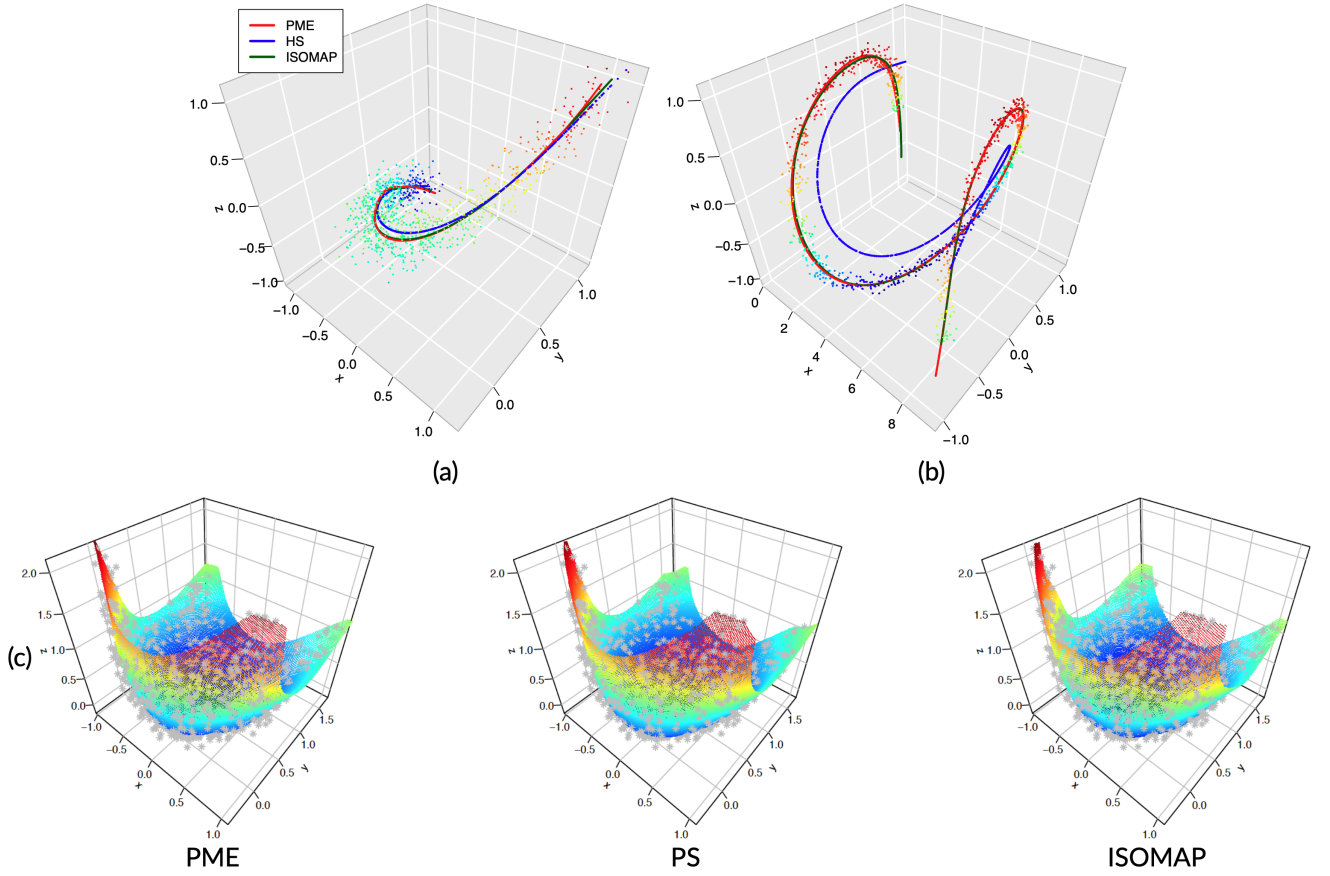


Figure 1.6: In each case, X are generated as follows: (a) using the same approach as Case I in Figure 1.4; (b) $(X|T = t) \sim N_3(t, 0.05I_{3 \times 3})$, $T = (\tau, \cos \tau, \sin \tau)^T$, $\tau \sim \text{Unif}(\pi/2, 6\pi)$; the three lower panels share the same data (gray), where $(X|T = t) \sim N_3(t, 0.05I_{3 \times 3})$, $\tau_1, \tau_2 \sim_{iid} \text{Unif}(-1, 1)$, and $T = (\tau_1, \frac{1}{2}(\tau_2 + \sqrt{3}(\tau_1^2 + \tau_2^2)), \frac{1}{2}(\tau_1^2 + \tau_2^2 - \sqrt{3}))^T$.

size of simulated data sets, applying ISOMAP becomes infeasible.

1.8 Interior Identification

In this Section, we propose an algorithm to identify the interiors of circle-like curves ($d = 1, D = 2$) and cylinder/ball-like surfaces ($d = 2, D = 3$). Examples of such curves and surfaces are presented in Figure 1.7. In many applications, the target is not the surface of an object, but its interior. For example, radiation therapists may be interested in identifying the interior of a tumor, which contains malignant cells. We propose an interior identification method based on PME.

Let $\underline{M}^d$ denote a circle-like curve ($d = 1$) or cylinder/ball-like surface ($d = 2$) contained in a domain $\mathcal{E} \subset \mathbb{R}^D$, e.g., the punched sphere in Figure 1.7 (b) is contained in a cube. The main idea of the interior identification approach is as follows: **Stage 1**, we decompose $\mathcal{E}$ into several potentially overlapping subsets, i.e., $\mathcal{E} = \bigcup_{s=1}^S E[s]$, where $E[s]$ are subsets of $\mathcal{E}$; **Stage 2**, we identify the interior of each piece $\underline{M}^d \cap E[s]$; **Stage 3**, we “glue” the piecewise interior identification results of all $\underline{M}^d \cap E[s]$ using the *10-nearest neighborhood classifier* and obtain

the interior estimation of $\underline{M}^d = \bigcup_{s=1}^S (\underline{M}^d \cap E[s])$. In Stage 2, we assume that, for each index s , there exists an $f_{\underline{s}} : \mathbb{R}^d \rightarrow \mathbb{R}^D$ such that $\underline{M}^d \cap E[s] = M_{f_{\underline{s}}}^d \cap E[s]$, where the manifold $M_{f_{\underline{s}}}^d$ is defined by Definition 1.2.1 and estimated using PME.

We first propose the interior identification approach for each piece $M_f^d \cap E$, where $f : \mathbb{R}^d \rightarrow \mathbb{R}^D$ and E is a sub-domain of $\mathcal{E}$. Let $\vec{n}(t)$ denote a normal vector of M_f^d at point $f(t)$. For example,

$$\begin{aligned} \vec{n}(t) &= \left(-\frac{df_2}{dt}(t), \frac{df_1}{dt}(t) \right)^T \quad \text{when } d=1 \text{ and } D=2, \text{ and} \\ \vec{n}(t) &= \left(\frac{\partial f_2}{\partial t_1} \frac{\partial f_3}{\partial t_2} - \frac{\partial f_3}{\partial t_1} \frac{\partial f_2}{\partial t_2}, \frac{\partial f_3}{\partial t_1} \frac{\partial f_1}{\partial t_2} - \frac{\partial f_1}{\partial t_1} \frac{\partial f_3}{\partial t_2}, \frac{\partial f_1}{\partial t_1} \frac{\partial f_2}{\partial t_2} - \frac{\partial f_2}{\partial t_1} \frac{\partial f_1}{\partial t_2} \right)^T \quad \text{when } d=2 \text{ and } D=3. \end{aligned}$$

Computing $\vec{n}(t)$ is possible since we have the analytic formula (1.6.2). For a fixed $\xi \in \mathbb{R}^D$,

$$Orit(\xi, f) = \text{sgn} \left\{ [f(\pi_f(\xi)) - \xi]^T \vec{n}(\pi_f(\xi)) \right\}$$

is called the *orientation* of ξ with respect to f , where $\text{sgn}(\cdot)$ is the sign function. Let c^* be a predetermined point indicating the interior side of $M_f^d \cap E$. It is called the *reference point*. Then all the points in E sharing the same orientation with c^* are identified as interior points, i.e., the interior part of $M_f^d \cap E$ is estimated by $\mathcal{I}(f, c^*) \cap E$, where $\mathcal{I}(f, c^*) = \{\xi \in \mathbb{R}^D : Orit(\xi, f) \times Orit(c^*, f) > 0\}$. A geometric illustration is presented in Figure 1.7 (a).

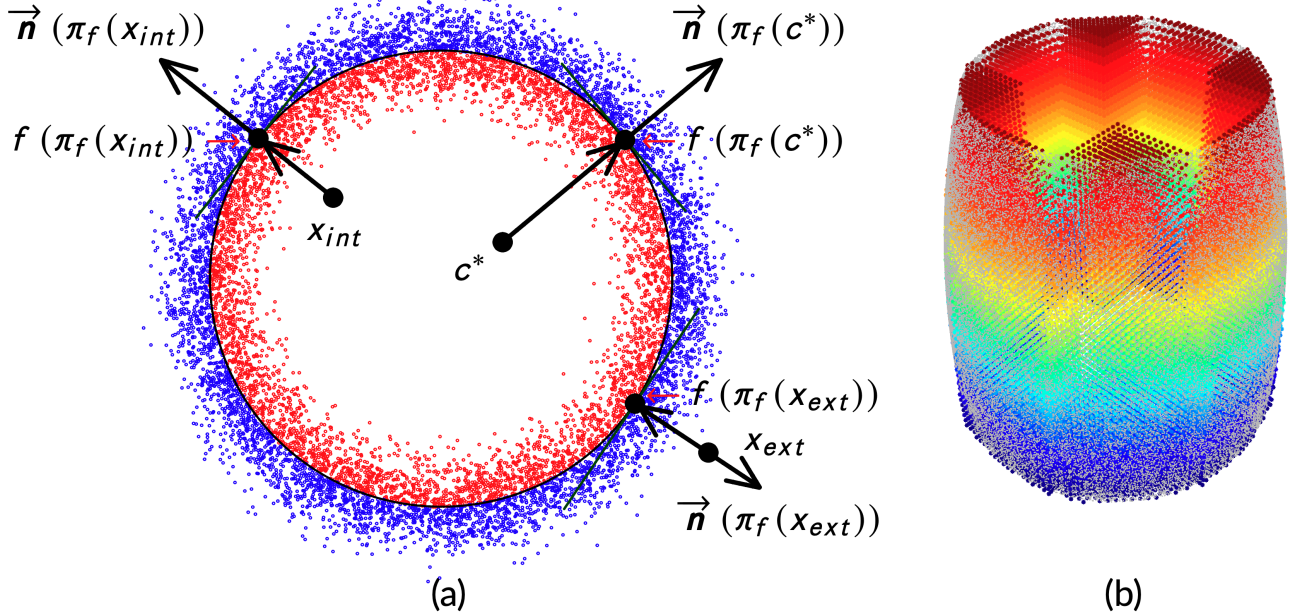


Figure 1.7: (a) An illustration of $\mathbf{n}(\pi_f(\cdot))$, $f(\pi_f(\cdot))$, and the reference point c^* . (b) A simulation example, where 10000 data points are from $(\sin \tau_1 \times \cos \tau_2, \sin \tau_1 \times \sin \tau_2, \cos \tau_1)^T$, with $\tau_1 \sim \text{Unif}(\pi/4, 3\pi/4)$, $\tau_2 \sim \text{Unif}(0, 2\pi)$. There is no 3-dimensional corrupting noise in these data. The colored points indicate the ξ_j identified as interior of the punched sphere. To illustrate the boundaries of the cubes $E[k]$, we omit all interior ξ_j outside of $\mathcal{E} = \bigcup_{k=1}^8 E[k]$.

Secondly, we explain the interior identification approach for the entire $\underline{M}^d$ by an example - fitting the $I = 10000$

data points $\{x_i\}_{i=1}^I$ in Figure 1.7 (b) (gray points). These data points are simulated from a punched sphere, which is the $\underline{M}^d$ of interest. The reference point $c^* = (0, 0, 0)^T$ is the centroid of the sphere. The points to be identified are grid-points ξ_j , such as the colored points in Figure 1.7 (b). In this example, we show constructions of the domain $\mathcal{E}$ and its subsets $E[s]$. We identify the grid-points ξ_j interior of this punched sphere using the following procedure.

Step 1: For each 3-dimensional vector $x_i = (x_{i,1}, x_{i,2}, x_{i,3})^T$ where $x_{i,l}$ denotes the l^{th} component of x_i , let $(\phi_i, r_i)^T$ be the polar coordinate of the 2-dimensional vector $(x_{i,1}, x_{i,2})^T$ and ϕ_i be the corresponding angle component. Partition $\{x_i\}_{i=1}^I$ into 8 subsets by $\mathcal{Z}[k] = \{x_i : \frac{(k-1)\pi}{4} \leq \phi_i < \frac{k\pi}{4}\}$ for $k = 1, 2, \dots, 8$.

Step 2: Define the cubes $E[k] = \prod_{l=1}^3 \left[\inf_{x_i \in \mathcal{Z}[k]} x_{i,l}, \sup_{x_i \in \mathcal{Z}[k]} x_{i,l} \right]$ for all k . Then $\mathcal{E} = \bigcup_{k=1}^8 E[k]$ contains all x_i .

Step 3: Fit an f_1 to data in $\mathcal{Z}[8] \cup \mathcal{Z}[1]$ and an f_k to data in $\mathcal{Z}[k-1] \cup \mathcal{Z}[k]$ for all $k = 2, 3, \dots, 8$ using PME.

Step 4: For each k , define $x_k^* = \frac{1}{\#\mathcal{Z}[k]} (\sum_{x_i \in \mathcal{Z}[k]} x_i) \in \mathbb{R}^3$, where $\#\mathcal{Z}[k]$ denotes the number of elements in $\mathcal{Z}[k]$.

Step 5: All grid-points $\xi_j \notin \mathcal{E}$ are identified as exterior. For each $\xi_j \in \mathcal{E}$, compute $k = \arg \min_{k'=1,2,\dots,8} \{\|\xi_j - x_{k'}^*\|_{\mathbb{R}^3}\}$. Since both f_k and f_{k+1} fit data in $\mathcal{Z}[k]$, there are three possible scenarios:

- (i) $\xi_j \in \mathcal{I}(f_k, c^*) \cap \mathcal{I}(f_{k+1}, c^*)$, i.e., ξ_j is identified as interior by both f_k and f_{k+1} , then ξ_j is identified as interior and labeled by “int”;
- (ii) $\xi_j \notin \mathcal{I}(f_k, c^*) \cup \mathcal{I}(f_{k+1}, c^*)$, i.e., ξ_j is identified as exterior by both f_k and f_{k+1} , then ξ_j is identified as exterior and labeled by “ext”;
- (iii) ξ_j satisfies neither the previous two scenarios, then we identify ξ_j by applying 10-nearest neighborhood classifier to the labeled training set $\{(\xi_q, lab_q) : \xi_q \in E[k] \text{ and } \xi_q \text{ satisfies scenario (i) or (ii)}\}$, where $lab_q \in \{\text{“int”}, \text{“ext”}\}$ is the label of ξ_q .

The performance of this interior identification procedure is shown in Figure 1.7 (b). Since we know the true punched sphere generating data, the true interior/exterior labels of grid-points $\xi_j \in \mathcal{E} = \bigcup_{k=1}^8 E[k]$ with respect to this punched sphere are known. The identification error rate - the proportion of incorrectly estimated labels - is less than 0.1%. In the illustrative example in Figure 1.7, we automatically and evenly divide data x_i into eight subcollections. In general, depending on the shape of the observed data, we may need to divide data into more/fewer subcollections. Additionally, an uneven division might be suitable for some data sets. For example, we may conduct a finer division in a region containing a large number of data points than in a region containing only a few data points. Determining the number of subcollections and division precision in individual regions is left for future research. Additionally, future research may extend our proposed methods for identifying the interiors of a more general set of manifolds.

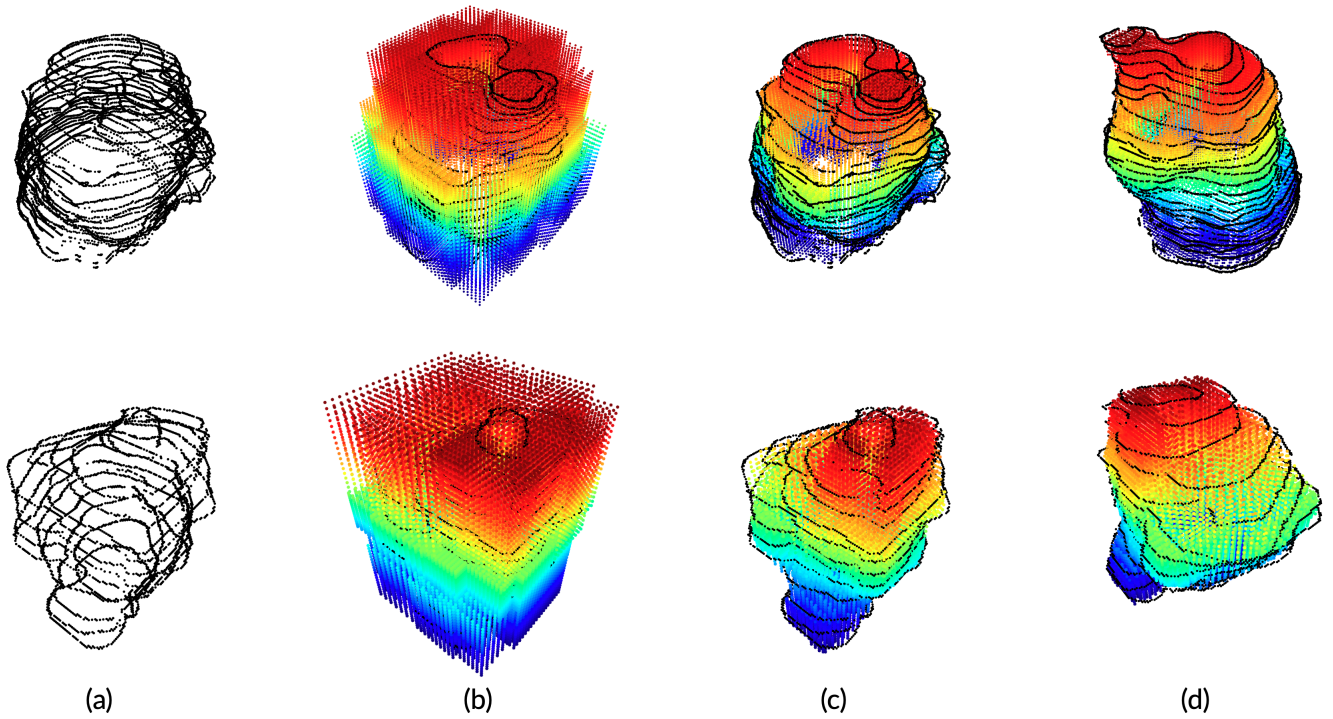


Figure 1.8: The black points in column (a) denote the CT data of two tumors. The colored points in column (b) denote the points to be identified. The colored points in column (c,d) denote the points identified as interior of the tumors. The last two columns show different angles of the tumors.

1.9 Analysis of Lung Cancer Tumor Data

In this section, we consider tumor surface estimation using computed tomography (CT) scans collected from patients with lung cancer and the identification of tumor interior in the context of radiation therapy. We analyzed two tumor data sets from a publicly available database collected for 422 patients with non-small cell lung cancer at the MAASTRO Clinic (Maastricht, The Netherlands) and available at <http://www.cancerimagingarchive.net/>. Spiral CT scans of the thoracic region with a $3mm$ slice thickness are obtained for each study participant. In addition, the masks of the tumor hand segmented by a radiologist are provided in the database. The hand segmentation result is a collection of voxels (3-dimensional counterparts of pixels) in 3-dimensional space marked by the radiologist as points on the surface of the tumor. The details on imaging parameters are available on the website, the references provided therein, and are not repeated in this section. The vertices of the tumors for the two participants are presented in Figure 1.8 (a). Given that we only have a collection of points on the tumor's surface, it is necessary to estimate the surface of the tumor fitted to the manually selected vertices on the surface of the tumor. In addition to estimating the tumor surface, it may be of interest to identify the interior area of the tumor. For example, in radiation therapy, ionizing radiation is used to control or kill cancer cells. To avoid harming healthy tissue with unnecessary doses of radiation, identifying the interior region of a tumor is important. Since the shape of the tumors is similar to a punched sphere, we apply the same procedures as in the example in

Section 1.8 to identify the interior part of these tumors.

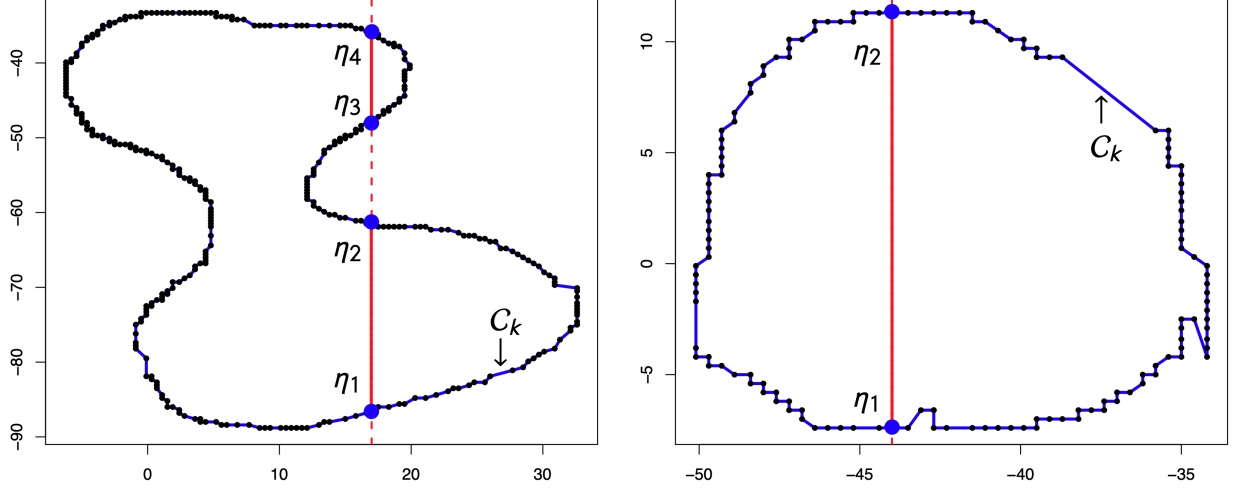


Figure 1.9: Left: a single slice from the CT data for one subject (presented in the upper panels of Figure 1.8). Right: a single slice from the CT data for another subject (presented in the lower panels of Figure 1.8).

The interior identification result is in Figure 1.8. Visually, the proposed method can properly identify interior points. We also use a simple approach to identify tumor interior points given the surface voxels and obtain a rough estimate of the validity of our proposed method. By its nature, the CT data is a collection of grid points in a 3D box $\mathcal{B} = [X_L, X_U] \times [Y_L, Y_U] \times [Z_L, Z_U]$ along with the intensities of all voxels in this grid. Let $\mathcal{X}^k = \{\xi_j^k\}_j$ be the collection of 3D data points in the k^{th} slice of the CT scan (all k here indicate the k^{th} slice). All the points in $\mathcal{X}^k$ share the same Z -coordinate $z^k \in [Z_L, Z_U]$. To identify the interior of the tumor in the rectangle $[X_L, X_U] \times [Y_L, Y_U] \times \{z^k\}$, e.g., the rectangles in Figure 1.9, we use a linear interpolation to connect consecutive points ξ_j^k and ξ_{j+1}^k . As a result, we obtain a piecewise linear and closed curve $\mathcal{C}_k$. The curve $\mathcal{C}_k$ (the blue curves in Figure 1.9) roughly indicates the boundary of the tumor in this slice. For any $x^k \in [X_L, X_U]$, let $\{\eta_1, \eta_2, \dots\}$ be the intersection of the line segment $\{x^k\} \times [Y_L, Y_U]$ and $\mathcal{C}_k$ (the blue dots in Figure 1.9). The points on line segments of the form

$$\bigcup_l \left\{ \lambda \eta_{2l-1} + (1 - \lambda) \eta_{2l} : \lambda \in [0, 1] \right\}$$

are identified as the tumor's interior. These line segments are the solid red ones in Figure 1.9.

Finally, for each of the $50 \times 50 \times \#\{z_k\}$ grid points ($\#\{z_k\}$ denotes the number of slices), we compare the labels given by the rough approximation approach in the preceding paragraph and the PME based interior identification method, respectively. For the two tumor data sets in Figure 1.8, 97.1% of the grid points are given the same labels by these two methods for subject 1 (top panel) and 95.4% for subject 2 (lower panel). Hence, we conclude that these two methods perform similarly in these two examples. However, the rough approximation method has major shortcomings. If the number of points identified by the hand segmentation is small, the line segments in the rough approach will result in a poor estimate of the tumor surface, leading to poor interior/exterior classification

performance. Additionally, any outlier surface voxels will potentially have significant negative effects on the rough classifier, while the PME approach is robust to the effects of outliers. Even though the rough approach has these limitations, we considered comparing it to our proposed approach as we have no gold standard classifier to illustrate the performance of our algorithm.

1.10 Conclusions and Discussion

In this paper, we propose a framework of principal manifolds for arbitrary intrinsic dimensions using Sobolev spaces. A Sobolev embedding theorem guarantees the regularity of principal manifolds. To reduce the computational cost and the effects of outliers, and to select model complexity, we propose the HDMDE algorithm motivated by [Eloyan and Ghosh \(2011\)](#). Future work may evaluate the performance of HDMDE as a density estimation technique in high-dimensional settings. Based on HDMDE, we develop the PME algorithm to estimate the newly proposed principal manifolds with intrinsic dimensions $d \leq 3$. We use simulations to compare PME to existing methods for scenarios with dimension pairs $(d = 1, D = 2)$, $(d = 1, D = 3)$, and $(d = 2, D = 3)$. These simulations illustrate that PME performs better than HS in many scenarios in the sense of minimizing MSD, ISOMAP is too computationally expensive compared to PME, and PME is not inferior to PS. However, PS is only defined for $d = 2$. Additionally, PS does not provide an explicit and simple formula of the map $f : \mathbb{R}^d \rightarrow \mathbb{R}^D$ defining estimated M_f^d , while we obtain such a formula using PME. We apply PME to radiation therapy by identifying the interiors of tumors, which are targets of ionizing radiation.

HS principal curves, which apply to data in Euclidean spaces, have been generalized to data on *Riemannian manifolds* (e.g., [Hauberg \(2015\)](#) and [Kim et al. \(2020\)](#)). In the meantime, PCA has been generalized to geodesics in *Wasserstein spaces* (e.g., [Boissard et al. \(2015\)](#) and [Seguy and Cuturi \(2015\)](#)). Future work may investigate the generalization of our proposed framework, for which the HS principal curve algorithm and PCA are special cases, to counterparts on Riemannian manifolds and in Wasserstein spaces. [Kirov and Slepčev \(2017\)](#) introduced a *multiple penalized principal curve* framework permitting a fitted 1-dimensional structure to consist of several disconnected curves. Our proposed framework may be generalized to fit a high-dimensional structure consisting of several disconnected components.

1.11 Appendix

Proof of Theorem 1.2.1: Since $\|f(t)\|_{\mathbb{R}^D} \rightarrow \infty$ as $\|t\|_{\mathbb{R}^d} \rightarrow \infty$, there exists $M > 0$ such that $\|x - f(t)\|_{\mathbb{R}^D} > 1 + \text{dist}(x, f)$ if $\|t\|_{\mathbb{R}^d} > M$. Then $\text{dist}(x, f) = \inf\{\|x - f(t)\|_{\mathbb{R}^D} : t \in \overline{B_d}(0, M)\}$, where $\overline{B_d}(0, M) = \{t \in \mathbb{R}^d : \|t\|_{\mathbb{R}^d} \leq M\}$. The compactness of $\overline{B_d}(0, M)$ implies that there exists $t^* \in \overline{B_d}(0, M)$ so that $\text{dist}(x, f) = \|x - f(t^*)\|_{\mathbb{R}^D}$, then $\mathcal{A}_f(x)$ is nonempty. We have

$$\mathcal{A}_f(x) = \{t \in \overline{B_d}(0, M) : \|x - f(t)\|_{\mathbb{R}^D} \leq \text{dist}(x, f)\} = \overline{B_d}(0, M) \cap f^{-1}(\overline{B_D}(x, \text{dist}(x, f))),$$

where $\overline{B_D}(x, \text{dist}(x, f)) = \{x' \in \mathbb{R}^D : \|x - x'\|_{\mathbb{R}^D} \leq \text{dist}(x, f)\}$ is compact, and $f^{-1}(\overline{B_D}(x, \text{dist}(x, f)))$ is closed as f is continuous. Since $\overline{B_d}(0, M)$ is bounded and closed, $\mathcal{A}_f(x)$ is compact. $\square$

Proof of Theorem 1.2.2 (i) can be proved following the method used to prove Theorem 4.1 in [Hastie \(1984\)](#). (ii) $\phi(x) = \text{dist}(x, f)$, then $\phi \in C(B)$ and $m^* = \max_{x \in B} \phi(x) < \infty$. Let $\zeta(t) = \inf_{x \in B} \|x - f(t)\|_{\mathbb{R}^D}$, then $\zeta(t) \geq \|f(t)\|_{\mathbb{R}^D} - \sup_{x \in B} \|x\|_{\mathbb{R}^D} \rightarrow \infty$ as $\|t\|_{\mathbb{R}^d} \rightarrow \infty$. $\exists M > 0$ so that $\zeta(t) > m^*$ if $\|t\|_{\mathbb{R}^d} \geq M$. For all $x \in B$, $\text{dist}(x, f) \leq m^* < \zeta(t) \leq \|x - f(t)\|_{\mathbb{R}^D}$ if $\|t\|_{\mathbb{R}^d} \geq M$, which implies $\{t : \|t\|_{\mathbb{R}^d} \geq M\} \cap [\bigcup_{x \in B} \mathcal{A}_f(x)] = \emptyset$. Then $\{\pi_f(x) : x \in B\} \subset \{t : \|t\|_{\mathbb{R}^d} < M\}$. (iii) Let $\Pi = f \circ \pi_f : \mathbb{R}^D \rightarrow M_f^d$. Since f has no ambiguity point on U , Theorem 1.3 of [Dudek and Holly \(1994\)](#) implies that Π is continuous on U . Since f^{-1} is continuous, $\pi_f = f^{-1} \circ \Pi$ is continuous on U . $\square$

Proof of Theorem 1.3.2: That $\mathcal{K}_{\infty, \mathbb{P}}(f) < \infty$ only if $\|\nabla^{\otimes 2} f\|_{L^2(\mathbb{R}^d)} = 0$ implies the generalized (not classical) derivatives $\frac{\partial^2 f_l}{\partial t_i \partial t_j} = 0$, for $1 \leq i, j \leq d$ and $l = 1, 2, \dots, D$, almost everywhere. From Lemma 1.3.1 and the Corollary 3.32 of [Adams and Fournier \(2003\)](#), f equals an affine function almost everywhere. The continuity of f implies that f equals this affine function exactly everywhere. Then $f(\pi_f(X))$ is the projection of X to some hyperplane. Therefore, $\inf_{f \in \mathcal{F}(\mathbb{P})} \mathcal{K}_{\infty, \mathbb{P}}(f) = \inf_{\mathbf{C} \in \mathcal{P}, a \in \mathbb{R}^D} \mathbb{E} \|X - (\mathbf{I} - \mathbf{C})a - \mathbf{C}X\|_{\mathbb{R}^D}^2$, where $\mathcal{P} = \{\mathbf{C} \in \mathbb{R}^{D \times D} : \mathbf{C}^2 = \mathbf{C}^T = \mathbf{C}, \text{rank}(\mathbf{C}) = d\}$ is the collection of projection matrices of rank d . Then $\inf_{f \in \mathcal{F}(\mathbb{P})} \mathcal{K}_{\infty, \mathbb{P}}(f)$ is equal to $\inf\{\mathbb{E} \|(\mathbf{I} - \mathbf{C})(X - a)\|_{\mathbb{R}^D}^2 : \mathbf{C} \in \mathcal{P}, a \in \mathbb{R}^D\} = \inf\{\text{tr}[(\mathbf{I} - \mathbf{C})\mathbf{U}\mathbf{D}\mathbf{U}^T(\mathbf{I} - \mathbf{C})^T] + \|(\mathbf{I} - \mathbf{C})(\mathbb{E}X - a)\|_{\mathbb{R}^D}^2 : a \in \mathbb{R}^D, \mathbf{C} \in \mathcal{P}\}$, where $\mathbf{U} = (\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_D)$ and $\mathbf{D} = \text{diag}(e_1, e_2, \dots, e_D)$. The minimum is achieved by the minimizer $(\mathbf{C}^*, a^*)$, where $\mathbf{C}^*$ is the projection matrix to the subspace $\{\sum_{i=1}^d \alpha_i \mathbf{v}_i : \alpha_i \in \mathbb{R}^1\}$ and a^* satisfies $(\mathbf{I} - \mathbf{C}^*)(\mathbb{E}X - a^*) = 0$. Then the minimizer hyperplane is $\{(\mathbf{I} - \mathbf{C}^*)a^* + \mathbf{C}^*x : x \in \mathbb{R}^D\} = \{\mathbb{E}X + \sum_{i=1}^d \alpha_i \mathbf{v}_i : \alpha_i \in \mathbb{R}^1\}$. $\square$

Derivation of (1.5.2): Let $\{Z_i\}_{i=1}^I$ define independent latent random variables taking values in $\{1, 2, \dots, N\}$, such that $(X_i | \theta_N, Z_i = z_i) \sim \psi_{\hat{\sigma}_N}(x_i - \mu_{z_i, N}) dx_i$ and $(Z_i | \theta_N) \sim \theta_{z_i, N}(\sum_{j=1}^N \delta_j(dz_i))$ for $i = 1, 2, \dots, I$. In other words, the latent variable Z_i indicates the class membership of the i^{th} observation in the mixture. Then we have $(X_i, Z_i) | \theta_N \sim \theta_{z_i, N} \psi_{\hat{\sigma}_N}(x_i - \mu_{z_i, N}) [dx_i \times \sum_{j=1}^N \delta_j(dz_i)]$ and $\mathbb{P}(Z_i = j | \theta_N, X_i = x_i) = w_{ij}(\theta_N)$. The complete likelihood of $\{(X_i, Z_i)\}_{i=1}^I$ with respect to the product measure $\prod_{i=1}^I \{dx_i \times \sum_{j=1}^N \delta_j(dz_i)\}$ is $L_C(\theta_N | x, z) =$

$\prod_{i=1}^I \theta_{z_i, N} \times \psi_{\sigma_N}(x_i - \mu_{z_i, N})$. For a fixed $\theta_N^{(k)} \in \Theta_N$, in the E-step of EM algorithm, we construct

$$\begin{aligned} Q(\theta_N | \theta_N^{(k)}) &= \mathbb{E} \left(\log L_C(\theta_N | X, Z) \mid X = x, \theta_N^{(k)} \right) \\ &= \sum_{i=1}^I \sum_{j=1}^N \left\{ w_{ij}(\theta_N^{(k)}) \log(\psi_{\sigma_N}(x_i - \mu_{z_i, N})) + w_{ij}(\theta_N^{(k)}) \log \theta_{j, N} \right\}. \end{aligned}$$

We implement the constraints $\int_{\mathbb{R}^D} x p(x | \theta_N) dx = \bar{x}$ and $\sum_{j=1}^N \theta_{j, N} = 1$ and obtain the Lagrangian

$$Q_\rho(\theta_N | \theta_N^{(k)}) = Q(\theta_N | \theta_N^{(k)}) + \rho_1 \left(1 - \sum_{j=1}^N \theta_{j, N} \right) + \rho_2^T \left(\bar{x} - \sum_{j=1}^N \theta_{j, N} \mu_{j, N} \right)$$

for $\rho_1 \in \mathbb{R}^1, \rho_2 \in \mathbb{R}^D$. Taking derivatives of $Q_\rho(\theta_N | \theta_N^{(k)})$, we obtain

$$\begin{aligned} \frac{\partial Q_\rho}{\partial \theta_{j, N}} &= \frac{1}{\theta_{j, N}} \sum_{i=1}^I w_{ij}(\theta_N^{(k)}) - \rho_1 - \rho_2^T \mu_{j, N} = 0 \text{ for all } j \text{ and} \\ \frac{\partial Q_\rho}{\partial \rho_1} &= 1 - \sum_{j=1}^N \theta_{j, N} = 0, \quad \frac{\partial Q_\rho}{\partial \rho_2} = \bar{x} - \sum_{j=1}^N \theta_{j, N} \mu_{j, N} = 0. \end{aligned}$$

The resulting equations for estimating $\theta_{j, N}$ are

$$\theta_{j, N} = \frac{\sum_{i=1}^I w_{ij}(\theta_N^{(k)})}{\rho_1 + \rho_2^T \mu_{j, N}}, \quad \sum_{j=1}^N \left(\frac{\sum_{i=1}^I w_{ij}(\theta_N^{(k)})}{\rho_1 + \rho_2^T \mu_{j, N}} \right) = 1, \quad \sum_{j=1}^N \left(\frac{\sum_{i=1}^I w_{ij}(\theta_N^{(k)})}{\rho_1 + \rho_2^T \mu_{j, N}} \right) \mu_{j, N} = \bar{x}$$

for all j , whose solution is given by (1.5.2). $\square$

Proof of Theorem 1.5.1: Define the weights $\theta_{j, N} = \int_{A_{j, N}} p(\mu) d\mu$, then $\int_{\mathbb{R}^D} p(\mu) d\mu = 1$ implies that θ_N is in the probability simplex Θ_N . By *Minkowski's inequality* (Theorem 2.9 of Adams and Fournier (2003)), we have

$$\begin{aligned} \|p_N(\cdot | \theta_N) - p\|_{L^q(\mathbb{R}^D)} &\leq \left(\sum_{j=1}^N \int_{A_{j, N}} \|\psi_{\sigma_N}(\cdot - (\mu_{j, N} - \mu)) - \psi_{\sigma_N}\|_{L^q(\mathbb{R}^D)} p(\mu) d\mu \right) + \|\psi_{\sigma_N} * p - p\|_{L^q(\mathbb{R}^D)} \\ &=: I_N + II_N. \end{aligned}$$

Since μ and $\mu_{j, N}$ are in $A_{j, N}$ and $\text{diam}(A_{j, N}) \leq d_N$, we have $I_N \leq \sup\{\|\psi_{\sigma_N}(\cdot - y) - \psi_{\sigma_N}\|_{L^q(\mathbb{R}^D)} : \|y\|_{\mathbb{R}^D} \leq d_N\} \rightarrow 0$ as $N \rightarrow \infty$. Applying Minkowski's inequality again, we have $II_N \leq \int_{\mathbb{R}^D} \|p(\cdot - \sigma_N \mu) - p\|_{L^q(\mathbb{R}^D)} \psi(\mu) d\mu$. Then the continuity of translations $p \mapsto p(\cdot - y)$ with respect to L^q -topology and the dominant convergence theorem imply $\lim_{N \rightarrow \infty} II_N = 0$. $\square$

Proof of Theorem 1.5.3: Since p and ψ are compactly supported, $\{\mu_{j, N}\}_{j=1}^N \subset \text{supp}(p)$ for all N , and $\lim_{M \rightarrow \infty} \sigma_N = 0$, there exists a compact B containing the support of p , $\psi_{\sigma_N}(\cdot - \mu_{j, N})$, and $p_N(\cdot | \theta_N) = \sum_{j=1}^N \theta_{j, N} \psi_{\sigma_N}(\cdot -$

$\mu_{j,N}$) for all N , and f has no ambiguity point in B . Then $|\mathcal{K}_{\lambda,P_N}(f) - \mathcal{K}_{\lambda,p}(f)| \leq H_N + I_N$, where

$$\begin{aligned} I_N &= \|p_N(\cdot|\theta_N) - p\|_{L^1(\mathbb{R}^D)} \times \sup_{x \in B} \|x - f(\pi_f(x))\|_{\mathbb{R}^D}^2, \\ H_N &= \sum_{j=1}^N \left(\sup_{N': N' \geq j} \theta_{j,N'} \right) \times (H_{j,N}^* + H_{j,N}^{**}), \\ H_{j,N}^* &= \left| \int_B \|x - f(\pi_f(x))\|_{\mathbb{R}^D}^2 [\psi_{\sigma_N}(x - \mu_{j,N}) dx - \delta_{\mu_j}(dx)] \right| \leq 2 \times \sup_{x \in B} \|x - f(\pi_f(x))\|_{\mathbb{R}^D}^2, \\ H_{j,N}^{**} &= \left| \int_B \|x - f(\pi_f(x))\|_{\mathbb{R}^D}^2 [\delta_{\mu_{j,N}}(dx) - \delta_{\mu_j}(dx)] \right| \leq 2 \times \sup_{x \in B} \|x - f(\pi_f(x))\|_{\mathbb{R}^D}^2, \quad \text{for all } N. \end{aligned}$$

$p_N(\cdot|\theta_N) \rightarrow p$ in L^1 implies $\lim_{N \rightarrow \infty} I_N = 0$. One can show $\lim_{N \rightarrow \infty} \mathcal{F}(\psi_{\sigma_N}(\cdot - \mu_{j,N})) = \lim_{N \rightarrow \infty} \mathcal{F}(\delta_{\mu_{j,N}}) = \mathcal{F}(\delta_{\mu_j})$, where $\mathcal{F}$ denotes Fourier transform. Since the Fourier transform of a probability is the characteristic function of this probability, *Levy continuity theorem* implies that the probability measure $\psi_{\sigma_N}(\cdot - \mu_{j,N}) dx$ converges to δ_{μ_j} weakly and $\delta_{\mu_{j,N}}$ converges to δ_{μ_j} weakly as $N \rightarrow \infty$. Theorem 1.2.2 implies the continuity of $\|x - f(\pi_f(x))\|_{\mathbb{R}^D}^2$ in B , and *Portmanteau theorem* implies $H_{j,N}^*, H_{j,N}^{**} \rightarrow 0$ as $N \rightarrow \infty$ for all j . Then *dominated convergence theorem* implies $\lim_{N \rightarrow \infty} H_N = 0$. $\square$

Lemma 1.11.1. (Theorem 4 bis of [Duchon \(1977\)](#)) Suppose $d \leq 3$. Let $\mathcal{C}$ be a finite subset of $\mathbb{R}^d$ such that every polynomial in $\text{Poly}_1[t]$ is uniquely determined by its values on $\mathcal{C}$. Then there exists exactly one function of the form $\sigma(t) = \sum_{c \in \mathcal{C}} s_c \eta_{4-d}(t - c) + p(t)$ taking prescribed values on $\mathcal{C}$, where $p \in \text{Poly}_1[t]$ and $\sum_{c \in \mathcal{C}} s_c q(c) = 0$ for all $q \in \text{Poly}_1[t]$. Moreover, if γ is another function taking the same prescribed values on $\mathcal{C}$, one has $\|\nabla^{\otimes 2} \sigma\|_{L^2} \leq \|\nabla^{\otimes 2} \gamma\|_{L^2}$.

Proof of Theorem 1.6.1: Lemma 1.11.1 implies that $g^* = \arg \min_{f \in \nabla^{-\otimes 2} L^2} \mathcal{K}_{\lambda, \widehat{Q}_N}(f, f_{(n)})$ is of the form (1.6.2). Theorem 1.3.1, $d \leq 3$, and the form of (1.6.2) imply $g^* \in C_\infty \cap \nabla^{-\otimes 2} L^2$. Hence,

$$g^* = \arg \min_{f \in C_\infty \cap \nabla^{-\otimes 2} L^2} \mathcal{K}_{\lambda, \widehat{Q}_N}(f, f_{(n)}) = f_{(n+1)}.$$

$\square$

Chapter 2

Population-level Task-evoked Functional Connectivity

2.1 Introduction

Functional magnetic resonance imaging (fMRI) is a non-invasive brain imaging technique used to estimate both brain regional activity and interactions between brain regions due to its ability to detect *neuronal activity* in the entire brain simultaneously. In fMRI studies, brain signals are measured on 3D volume elements (*voxels*) during a certain period of time, and each signal is observed at discrete time points. For example, if an fMRI study lasts for 20 minutes and data are collected at every 2 seconds, that is, the time between successive brain scans referred to as *repetition time* (TR, denoted as Δ hereafter) is 2 seconds, then we obtain observations at 900 time points. It is often of interest to investigate neuronal activity. However, since neuronal activity occurs in milliseconds, it is impossible to directly observe it using fMRI technology. It is implicitly captured by *blood-oxygenation level-dependent* (BOLD) signals. When neuronal activity in a brain area occurs, it is followed by localized changes in metabolism, where the corresponding local oxygen consumption increases, and then oxygen-rich blood flows to this area. This process results in an increase in oxyhemoglobin and a decrease in deoxyhemoglobin. The BOLD signal value at each time

- A version of this chapter appeared in [Meng and Eloyan \(2021a\)](#).
- **Author contributions:** Kun Meng developed the theoretical framework of this chapter, proved the theorems, and conducted the simulation studies and the HCP data application. Ani Eloyan supervised this chapter. Kun Meng and Ani Eloyan wrote version [Meng and Eloyan \(2021a\)](#). We — Kun Meng and Ani Eloyan — want to thank Dr. Matthew T. Harrison for his comment on the interpretation of our proposed ptFC, and this comment improved the interpretation of our model.
- **Abbreviations:** BOLD, blood-oxygenation level-dependent; DC, dynamic connectivity; fMRI, functional magnetic resonance imaging; FC, functional connectivity; HCP, Human Connectome Project; HRF, hemodynamic response function; ptFC, task-evoked functional connectivity at the population level; ptFCE, ptFC estimation; ROI, region of interest; SLR, subject-wise linear regression; WSMZ, weakly stationary with mean zero.

point is the difference between the oxyhemoglobin and deoxyhemoglobin levels. The signal consisting of BOLD values across all time points measures the localized metabolic activity influenced by the local brain vasculature, and it indirectly measures the localized neuronal activity. fMRI captures BOLD signals while study subjects are either resting (resting-state fMRI) or performing tasks (task-fMRI).

Our brains are networks that can be expressed as graphs. Dependencies between activation in any two regions of the brain are referred to as *functional connectivity* (FC, [Cribben and Fiecas \(2016\)](#)). These dependencies have been widely studied, and several approaches are proposed for estimation of FC, especially at rest. Task-evoked FC is fundamentally different from resting-state FC. [Lynch et al. \(2018\)](#) show that existing estimates of task-evoked FC cannot explain differences between FC at rest and during a task. While brain FC in an individual subject is often of interest, referred to as subject-level FC, population-level estimates of FC have been proposed in the literature and used for obtaining more robust subject-level FC estimates (e.g., see [Mejia et al. \(2018\)](#)). To overcome the limitations of existing task-evoked FC estimation approaches when investigating population-level FC, we propose a novel definition of *population-level task-evoked FC* (ptFC). This definition is based on a correlation of subject-specific random effects capturing task-evoked neuronal activity. Importantly, our proposed ptFC framework takes into account the complex biological processes of the brain response to task stimuli that many existing estimators ignore.

We are interested in modeling neuronal activity and BOLD signals of a set of *subjects* ω from a population of interest. At each time point t , we denote by $Y_k(\omega; t)$ the BOLD signal value at the k^{th} node of ω 's brain. The word “node” herein refers to either a voxel or a *region of interest* (ROI). Aggregated BOLD signals at a macro-area level are often of interest. That is, for each subject, BOLD signals are spatially averaged within pre-selected ROI and the resulting ROI-specific signals are analyzed. The BOLD signals for subject ω are represented by a K -vector-valued function $\{\mathbf{Y}(\omega; t) = (Y_1(\omega; t), \dots, Y_K(\omega; t))^T \mid t \in \mathcal{T}\}$ on $\mathcal{T}$, where $\mathcal{T}$ is the collection of time indices and K is the number of nodes. Furthermore, we assume that $\mathcal{T}$ is a compact subset of $\mathbb{R}$. In this paper, we model three types of signals in task-fMRI studies: (i) *stimulus signals* denoted by $N(t)$ representing the experimental designs of tasks, (ii) *task-evoked neuronal activity signals* at the k^{th} node $\Phi_k[\omega; N(t)]$ stemming solely from the task stimulus $N(t)$, where the “stimulus-to-activity” maps $\Phi_k[\omega; \bullet] : N(t) \mapsto \Phi_k[\omega; N(t)]$ depend on subjects ω , and (iii) observed BOLD signals $Y_k(\omega; t) = \Psi_k\{\omega; \Phi_k[\omega; N(t)]\}$ that are associated with the task-evoked neuronal activity $\Phi_k[\omega; N(t)]$, where the “activity-to-BOLD” maps $\Psi_k\{\omega; \bullet\} : \Phi_k[\omega; N(t)] \mapsto Y_k(\omega; t)$ are ω -dependent. In (ii) and (iii), map $\Phi_k[\omega; \bullet]$ characterizes how neurons at the k^{th} node of ω react to stimulus $N(t)$, and map $\Psi_k\{\omega; \bullet\}$ describes how neuronal activity $\Phi_k[\omega; N(t)]$ induces BOLD signal $Y_k(\omega; t)$. We are interested in modeling $\Phi_k[\omega; N(t)]$. However, only BOLD signals $Y_k(\omega; t)$ are observable. The goal of many fMRI studies, including our work, is to recover $\Phi_k[\omega; \bullet]$ by analyzing $Y_k(\omega; t)$. In [Section 2.2](#), we provide nonparametric models for $\Phi_k[\omega; \bullet]$ and $\Psi_k\{\omega; \bullet\}$.

Theoretically, time t is continuous and $\mathcal{T} = [0, t^*]$, where $t^* < \infty$ denotes the end of experiment. In applications,

we obtain data only at discrete and finite time points in $\mathcal{T} = \{\tau\Delta\}_{\tau=0}^T$, where Δ is the predetermined TR and T indicates that BOLD signals are observed at $T + 1$ time points. A BOLD signal at the k^{th} node of subject ω in task-fMRI consists of three components: (i) $P_k(\omega; t)$ denotes the one evoked solely by the experimental task $N(t)$, (ii) $Q_k(\omega; t)$ denotes the one stemming from spontaneous brain activity, e.g., the activity coordinating respiration and heartbeat, and the neuronal activity responding to stimuli that we are not of interest, and (iii) random error $\epsilon_k(\omega; t)$. We assume that these components have an additive structure, and observed BOLD signals $Y_k(\omega; t)$ are of the following form.

$$Y_k(\omega; t) = P_k(\omega; t) + Q_k(\omega; t) + \epsilon_k(\omega; t), \quad k = 1, 2, \dots, K, \quad t \in \mathcal{T}, \quad (2.1.1)$$

where $P_k(\omega; t)$ are called *task-evoked terms*. $P_k(\omega; t)$ are of primary interest and identifiable under some probabilistic conditions. The proof of identifiability of $P_k(\omega; t)$ is in Appendix C. Model (2.1.1) is equivalent to many existing approaches for modeling BOLD signals in task-fMRI (e.g., Joel et al. (2011), Zhang et al. (2013), and Warnick et al. (2018)). The relationship between (2.1.1) and these approaches is discussed in Appendix A. Similar to these existing methods, we assume no interaction between $P_k(\omega; t)$ and $Q_k(\omega; t)$ in (2.1.1). Inclusion of an interaction between these terms is discussed in Appendix H.

We implement our proposed methods to estimate functional connectivity in a study investigating the human motor system. To investigate motor FC, we use a cohort of subjects from a task-evoked fMRI study publicly available at the Human Connectome Project (HCP). The *block-design* motor task used in this study is adapted from experiments by Buckner et al. (2011) and Yeo et al. (2011), while details on the HCP implementation are in Barch et al. (2013). During the experiment, the subjects are asked to perform five tasks when presented with a cue: tap left/right fingers, squeeze left/right toes, and move their tongue. BOLD signals $Y_k(\omega; t)$ are collected from 308 participants, and each BOLD signal is obtained with $\Delta = 0.72$ (seconds). Each experiment lasts for about 204 seconds, i.e., $T = 283$. We focus on the brain functional connectivity for the task of *squeezing right toes*. The onsets of these task blocks vary across subjects. Nevertheless, the corresponding onsets of any two participants differ in less than 0.1 seconds ($\leq \Delta$). Therefore, we assume that all subjects share the same stimulus signal $N(t) = \mathbf{1}_{[86.5, 98.5)}(t) + \mathbf{1}_{[162, 174)}(t)$.

In statistical analysis of task-fMRI data $\{\mathbf{Y}(\omega; t)\}_{t \in \mathcal{T}}$ in (2.1.1), two topics are primarily of interest: (i) identification of nodes presenting task-evoked neuronal activity, i.e., the indices k such that $P_k(\omega; t) \neq 0$, and (ii) detection of task-evoked associations between brain nodes. *General linear models* (GLM, Friston et al. (1994)) are commonly implemented to detect task-evoked nodes (Lindquist et al. (2008)). GLM is conducted at an individual node level and is not informative for investigating associations between nodes. Associations between nodes captured by BOLD signals characterize FC (Friston et al. (1993)). FC observed during a task experiment tends to be different from that observed in a resting-state experiment (Lowe et al. (2000)). The difference between task-fMRI and resting-

state fMRI is presented by the task-evoked terms $P_k(\omega; t)$. In this paper, we investigate the associations between $\{P_k(\omega; t)\}_{k=1}^K$, instead of those between BOLD signals $\{Y_k(\omega; t)\}_{k=1}^K$, corresponding to FC stemming solely from the task of interest $N(t)$.

A considerable amount of work has been done to define and estimate FC. For example, [Friston et al. \(1993\)](#) defines FC as the temporal Pearson correlation between a pair of BOLD signals across time. Following this approach, one may measure task-evoked FC between $P_k(\omega; t)$ and $P_l(\omega; t)$ by the absolute value of the Pearson correlation as follows

$$|corr(\omega; P_k, P_l)| := \left| \frac{\int_{\mathcal{T}} P_k^*(\omega; t) \times P_l^*(\omega; t) \mu(dt)}{\sqrt{\int_{\mathcal{T}} |P_k^*(\omega; t)|^2 \mu(dt) \times \int_{\mathcal{T}} |P_l^*(\omega; t)|^2 \mu(dt)}} \right|, \quad (2.1.2)$$

where $P_{k'}^*(\omega; t) = P_{k'}(\omega; t) - \frac{1}{\mu(\mathcal{T})} \int_{\mathcal{T}} P_{k'}(\omega; s) \mu(ds)$ for $k' \in \{k, l\}$; if $\mathcal{T} = [0, T^*]$, then $\mu(dt) = dt$; if $\mathcal{T} = \{\tau\Delta\}_{\tau=0}^T$, then $\mu(dt)$ is the counting measure $\sum_{\tau \in \mathbb{Z}} \delta_{\tau\Delta}(dt)$, where $\delta_{\tau\Delta}$ is the point mass at $\tau\Delta$. There are many other approaches to defining and estimating FC, e.g., *coherence analysis* ([Müller et al. \(2001\)](#)) and *beta-series regression* ([Rissman et al. \(2004\)](#)) are widely used in FC studies. These approaches have several limitations discussed in Section 2.2.2 that are addressed by our proposed approach.

This paper is organized as follows. Section 2.2 proposes models for task-evoked neuronal activity and BOLD signals, respectively; based on these models, we propose a rigorous and interpretable definition of ptFC. Section 2.3 presents an algorithm for estimating ptFC. We compare the performance of our proposed algorithm with existing approaches using simulations in Section 2.4. In Section 2.5, we apply the proposed approach to estimate the ptFC during a motor task using the publicly available HCP data set. Section 2.6 concludes the paper.

2.2 Population-level Task-evoked Functional Connectivity (ptFC)

We first propose models for neuronal activity and BOLD signals at individual nodes. Using these models, we provide the definition of ptFC. Lastly, we list advantages of ptFC compared to existing approaches. Hereafter, Ω denotes a study population of interest, i.e., the collection of subjects of interest. Ω is discrete and finite. Let $\mathbb{P}$ define a study-dependent probability on Ω . Then $(\Omega, 2^\Omega, \mathbb{P})$ is a probability space, where 2^Ω is the power set of Ω . $\mathbb{E}$ denotes the expectation with respect to $\mathbb{P}$. In the HCP experiment described in Section 2.1, Ω represents the population of healthy adults. FMRI data are collected for each subject ω in a sample of size $n = 308$ drawn from Ω according to the underlying distribution $\mathbb{P}$. We estimate $\mathbb{P}$ by empirical distribution $\frac{1}{n} \sum_{\omega} \delta_{\omega}$. The expectation and correlation with respect to this empirical distribution are the sample average and sample correlation, respectively.

2.2.1 Definition of ptFC

Let $h_k(t)$ denote the *hemodynamic response function* (HRF, Lindquist et al. (2008)) at the k^{th} node corresponding to the task of interest $N(t)$ and shared by all subjects $\omega \in \Omega$. For each $\omega \in \Omega$, the GLM approach essentially models the task-evoked terms in (2.1.1) by

$$P_k(\omega; t) = \beta_k(\omega) \times (N * h_k)(t - t_{0,k}), \quad t \in \mathcal{T}, \quad \omega \in \Omega, \quad k = 1, 2, \dots, K. \quad (2.2.1)$$

where $\beta_k(\omega)$ is a GLM regression coefficient representing the subject-specific random effects, $t_{0,k}$ is the latency shared by all subjects as the reaction of the k^{th} node to $N(t)$ at time t is not instantaneous, and $*$ denotes convolution. Specifically, this is a GLM with response $\{P_k(\omega; t)\}_{t \in \mathcal{T}}$ and independent variable $\{(N * h_k)(t - t_{0,k})\}_{t \in \mathcal{T}}$. The error term of the GLM is absorbed by $\{\epsilon_k(\omega; t)\}_{t \in \mathcal{T}}$ in (2.1.1). Task-evoked terms $P_k(\omega; t)$ in (2.2.1) model the neuronal activity evoked by task $N(t)$, and the subject-specific $\beta_k(\omega)$ measures the magnitude of this component. Visualizations of (2.2.1) are in Web Figure 2.9. Theoretically, latency depends on subject ω . However, in most studies, the latency is much shorter than the corresponding TR Δ , and the difference between the latency times of any two subjects is negligible. Therefore, we model the latency as $t_{0,k}$ shared by all subjects. With model (2.2.1), we represent the model (2.1.1) for BOLD signals as follows.

$$Y_k(\omega; t) = \beta_k(\omega) \times (N * h_k)(t - t_{0,k}) + R_k(\omega; t), \quad t \in \mathcal{T}, \quad k = 1, \dots, K, \quad \omega \in \Omega, \quad (2.2.2)$$

where $R_k(\omega; t) = Q_k(\omega; t) + \epsilon_k(\omega; t)$ are referred to as *reference terms* for succinctness.

Task-evoked terms $P_k(\omega; t) = \beta_k(\omega) \times (N * h_k)(t - t_{0,k}) = [\{\beta_k(\omega) \times N(\cdot - t_{0,k})\} * h_k](t)$ correspond to neuronal activity evoked by $N(t)$, where h_k represent metabolism and vasculature and do not characterize neuronal activity. Therefore, we model neuronal activity signals $\Phi_k[\omega; N(t)]$ responding to $N(t)$ as $\beta_k(\omega) \times N(t - t_{0,k})$. Using model (2.2.2), the “stimulus-to-activity” map $\Phi_k[\omega; \bullet]$ and “activity-to-BOLD” map $\Psi_k\{\omega; \bullet\}$ in Section 2.1 are $\Psi_k\{\omega; \bullet\} : \Phi_k[\omega; N(t)] \mapsto \{\Phi_k[\omega; N(\cdot)] * h_k\}(t) + Q_k(\omega; t) + \epsilon_k(\omega; t) = Y_k(\omega; t)$ and

$$\Phi_k[\omega; \bullet] : N(t) \mapsto \beta_k(\omega) \times N(t - t_{0,k}). \quad (2.2.3)$$

FC is anticipated to characterize the mechanism of neuronal activity rather than metabolism or vasculature. Therefore, task-evoked FC is defined using task-evoked neuronal activity signals $\Phi_k[\omega; N(t)]$. In model (2.2.3) of $\Phi_k[\omega; N(t)]$, tasks $N(t)$ are fully determined by experimental designs. Additionally, latency $t_{0,k}$ can be viewed as a parameter of HRF h_k since task-evoked terms can be expressed as $\{[\beta_k(\omega) \times N] * [h_k(\cdot - t_{0,k})]\}(t)$. Therefore, task-evoked FC is expected to be determined by $\beta_k(\omega)$. Since $\beta_k(\omega)$, for $k = 1, \dots, K$, measure the magnitude of the neuronal activity evoked by task $N(t)$, nodes k and l are functionally connected at the population level during the task if one of the following scenarios holds: (i) for subjects with strong reaction to $N(t)$ at their k^{th} nodes, their

l^{th} nodes' reaction to $N(t)$ is strong as well and vice versa; (ii) for subjects with strong reaction to $N(t)$ at their k^{th} nodes, their l^{th} nodes' reaction to $N(t)$ is weak and vice versa, i.e., either positively or negatively correlated. Therefore, we make the following assumptions on the distribution of $\beta_k(\omega)$ across $(\Omega, \mathbb{P})$: (1) if there exists a functional connection evoked by $N(t)$ between nodes k and l , the corresponding $\beta_k(\omega)$ and $\beta_l(\omega)$ are approximately linearly associated; (2) each variance $\mathbb{V}\beta_k = \mathbb{E}(\beta_k^2) - (\mathbb{E}\beta_k)^2 > 0$, i.e., random variable $\beta_k(\omega)$ is not deterministic, where $\mathbb{E}\beta_k^\nu = \int_{\Omega} \{\beta_k(\omega)\}^\nu \mathbb{P}(d\omega)$ for $\nu = 1, 2$. Since $corr(\beta_k, \beta_l) := \frac{\mathbb{E}(\beta_k\beta_l) - (\mathbb{E}\beta_k)(\mathbb{E}\beta_l)}{\sqrt{\mathbb{V}\beta_k \times \mathbb{V}\beta_l}}$ measures the linear correlation between $\beta_k(\omega)$ and $\beta_l(\omega)$, we define ptFCs as follows.

Definition 2.2.1. For each subject $\omega \in \Omega$, the task-evoked BOLD signals $\{Y_k(\omega; t)\}_{k=1}^K$ are of form (2.2.2). The ptFC between the k^{th} and l^{th} nodes is defined as $|corr(\beta_k, \beta_l)|$.

An advantage of ptFC is its scale-invariance. Since $\beta_k(\omega)(N * h_k) = \frac{\beta_k(\omega)}{c_1 \times c_2} [(c_2 N) * (c_1 h_k)]$, the scale of $\beta_k(\omega)$ changes if the scale of h_k or $N(t)$ changes. But $|corr(\beta_k, \beta_l)|$ is invariant to the transform $\beta_{k'}(\omega) \mapsto c\beta_{k'}(\omega)$, for $k' \in \{k, l\}$ and any $c \neq 0$. Furthermore, using this scale-invariance, we show in Appendix C that ptFC is identifiable under some probabilistic conditions. Additionally, $|corr(\beta_k, \beta_l)|$ is invariant to the transform $\beta_{k'}(\omega) \mapsto \beta_{k'}(\omega) + c$ for $k' \in \{k, l\}$ and any $c \in \mathbb{R}$, hence we may assume $\mathbb{E}\beta_{k'} = 0$ for $k' \in \{k, l\}$ in Section 2.3.

Interpretation of $\beta_k(\omega)$: Each signal $\Phi_k[\omega; N(t)] = \beta_k(\omega) \times N(t - t_{0,k})$ describes the neuronal activity evoked by task $N(t)$, excluding effects of local vasculature or metabolism. If the k^{th} node of ω does not react to task $N(t)$, then $\beta_k(\omega) = 0$. For given $N(t)$ and h_k , the magnitude of $|\beta_k(\omega)|$ indicates the strength of the reaction in k^{th} node of ω to $N(t)$.

Interpretation of $|corr(\beta_k, \beta_l)|$: The pair $\{\beta_k(\omega), \beta_l(\omega)\}_{\omega \in \Omega}$ quantifies the magnitude of neuronal activity in response to $N(t)$ for nodes k and l . These two nodes have functional connectivity evoked by $N(t)$ if $\{\beta_k(\omega)\}_{\omega \in \Omega}$ and $\{\beta_l(\omega)\}_{\omega \in \Omega}$ are linearly associated across $(\Omega, \mathbb{P})$. More explicitly, a perfectly linear relationship between $\{\beta_k(\omega)\}_{\omega \in \Omega}$ and $\{\beta_l(\omega)\}_{\omega \in \Omega}$ implies strongest functional connectivity between nodes k and l evoked by $N(t)$. Finally, $|corr(\beta_k, \beta_l)|$ quantifies the strength of $N(t)$ -induced connectivity.

The proposed framework for estimation of population-level FC results in interpretable quantification of FC based on a more realistic model of task-evoked neuronal activity relationships between brain nodes compared with the existing methods discussed in Section 2.2.2. The Pearson correlation approach defines FC between two brain nodes by the correlation defined on time index space $\mathcal{T}$ (see equation (2.1.2)). In contrast, ptFC in Definition 2.2.1 is of a correlation form defined on the population space Ω . Although the $\mathcal{T}$ -correlation and Ω -correlation forms are different from the mathematical perspective, they model the same brain activity mechanism. Therefore, the correlations are comparable.

Recently, there is an emerging consensus in the literature suggesting that brain networks undergo temporal fluctuations corresponding to experimental tasks. *Dynamic connectivity* (DC) is a collection of approaches investigating these fluctuations. Although ptFC in Definition 2.2.1 does not change over time, the ptFC framework can be directly incorporated to obtain *sliding-window* based DC estimates (Hutchison et al. (2013)). Specifically, for

a preselected window of time, BOLD signals are extracted for that time interval, and ptFC is estimated using the algorithm proposed in this paper. Next, the window is shifted in time for a certain number of time points, and ptFC is estimated for BOLD signals within the time interval of the same length as in the first step for the shifted time interval. As a result, we obtain a sequence of ptFCs illustrating the dynamic structure of FC.

2.2.2 Advantages of ptFC

In this subsection, we present advantages of ptFC in Definition 2.2.1 compared to existing approaches. We first discuss the limitations of Pearson correlation approach (2.1.2) using model (2.2.1). Plugging (2.2.1) into (2.1.2), we obtain the following association between $P_k(\omega; t)$ and $P_l(\omega; t)$.

$$|corr(\omega; P_k, P_l)| = \left| \int_{\mathcal{T}} \phi_k(t) \times \phi_l(t) \mu(dt) \right| / \sqrt{\int_{\mathcal{T}} |\phi_k(t)|^2 \mu(dt) \times \int_{\mathcal{T}} |\phi_l(t)|^2 \mu(dt)}, \quad (2.2.4)$$

where $\phi_{k'}(t) = N * h_{k'}(t - t_{0,k'}) - \frac{1}{\mu(\mathcal{T})} \int_{\mathcal{T}} N * h_{k'}(s - t_{0,k'}) \mu(ds)$ for $k' \in \{k, l\}$. (2.2.4) reveals the following limitations of Pearson correlation approach (2.1.2).

Only nuisance parameters: Based on the reasoning in Section 2.2.1, task-evoked neuronal activity is modeled by $\beta_k(\omega)$. Hence, it is counterintuitive that (2.2.4) does not depend on $\beta_k(\omega)$ and depends only on the nuisance parameters $N(t)$ and $h_k(t)$ when considering neuronal activity. For example, if the k^{th} node does not react to $N(t)$, then the task-evoked term $\beta_k(\omega) \times (N * h_k)(t - t_{0,k})$ is expected to be zero, i.e., $\beta_k(\omega) = 0$, and there should be no task-evoked interaction between the k^{th} and other nodes. However, (2.2.4) can still be very large as the non-reaction information presented by $\beta_k(\omega) = 0$ vanishes. In contrast with (2.2.4), Definition 2.2.1 is based on $\{(\beta_k(\omega), \beta_l(\omega)) | \omega \in \Omega\}$.

Variation in latency: $t_{0,k}$ may vary across nodes, e.g., see Miezin et al. (2000) for investigation of the left and right visual and motor cortices and the corresponding difference between response onsets. If $t_{0,k} \neq t_{0,l}$, for example $t_{0,k} < t_{0,l}$, then it is likely that node k reacts to task $N(t)$ first, and then the neuronal activity at node k causes that at node l . Because of this potential causality represented by $t_{0,k} < t_{0,l}$, it is natural to expect that nodes k and l are likely functionally connected. However, measurement (2.2.4) can be very small if $t_{0,k} \neq t_{0,l}$ and may not reveal the true interaction between neuronal activity of two nodes. An example of this issue is illustrated in Appendix F. Since our proposed ptFCs do not involve latency, the variation in latency does not influence ptFCs in Definition 2.2.1.

Variation in HRFs: HRFs can heavily vary across brain nodes (Miezin et al. (2000)). Since task-evoked FC is not expected to depend on HRFs, variation of HRFs in different brain regions should not influence task-evoked FC. However, (2.2.4) can be small if $h_k \neq h_l$ as illustrated in Web Figure 2.9 (c). Given that ptFCs do not involve HRFs, they are invariant to the variation in HRFs across brain nodes.

Coherence analysis for FC is not influenced by any of the issues discussed above. In this approach, FC between

BOLD signals $Y_k(\omega; t)$ and $Y_l(\omega; t)$ is measured by the *coherence* evaluating the extent of the *linear time-invariant relationship* between these two signals via Fourier frequencies. However, there is no guarantee that two BOLD signals are linear time-invariant if the corresponding two nodes are functionally connected. Last but not least, the Pearson correlation and coherence analysis approaches are designed to measure FC evoked by all stimuli - both the tasks of interest and nuisance stimuli. Hence, they are not interpretable from the task-evoked FC viewpoint. On the contrary, beta-series regression and ptFCs are designed to measure the FC evoked by an experiment's specific task of interest. Additionally, our simulation studies in Section 2.4 show that beta-series regression performs worse than our proposed ptFC approach in many cases.

2.3 ptFC Estimation (ptFCE) Algorithm

Since BOLD signals are observed at discrete and finite time points, we assume $\mathcal{T} = \{\tau\Delta\}_{\tau=0}^T$ and $t = \tau\Delta$ for $\tau \in \{0, 1, \dots, T\}$. To estimate $|corr(\beta_k, \beta_l)|$, one needs the task-evoked terms $P_k(\omega; t)$. However, these terms are not observed in most applications. In Section 2.3.1, we propose an estimator defined as $\tilde{R}_k(\omega; t)$ for $R_k(\omega; t)$. As a result, we can obtain estimators for task-evoked terms as $\tilde{P}_k(\omega; t) := Y_k(\omega; t) - \tilde{R}_k(\omega; t)$. Suppose h_k and $t_{0,k}$ are known, then a straightforward approach to estimate $|corr(\beta_k, \beta_l)|$ is the subject-wise linear regression (SLR) as follows: For each fixed ω , we regress $\{\tilde{P}_k(\omega; \tau\Delta)\}_{\tau=0}^T$ on $\{(N * h_k)(\tau\Delta - t_{0,k})\}_{\tau=0}^T$ without intercept, and the ordinary least squares estimator for subject-dependent coefficient $\beta_k(\omega)$ is denoted as $\tilde{\beta}_k(\omega)$; then $|corr(\beta_k, \beta_l)|$ is estimated by the absolute value of the sample correlation between $\tilde{\beta}_k(\omega)$ and $\tilde{\beta}_l(\omega)$ across all ω . However, the SLR approach is not robust to the violation of assumptions in the model (2.2.2) and leads to worse performance of this approach compared to our proposed method as we present in Section 2.4 using simulations.

To obtain a better estimator of ptFC, in this section, we propose the ptFCE algorithm based on the Fourier transform and the AMUSE algorithm (Tong et al. (1991)). First, we provide notations for the Fourier transform. For any function $f : \mathcal{T} \rightarrow \mathbb{R}$, we extend f as follows to be a periodic function on $\{\tau'\Delta\}_{\tau' \in \mathbb{Z}}$.

$$f(\tau'\Delta) = f(\tau\Delta), \quad \tau' \equiv \tau \pmod{T+1} \text{ for } \tau = 0, 1, \dots, T \text{ and all } \tau' \in \mathbb{Z}. \quad (2.3.1)$$

Hereafter, all functions on $\mathcal{T}$ are implicitly extended using (2.3.1) to be periodic functions on $\{\tau'\Delta\}_{\tau' \in \mathbb{Z}}$. Convolution

$$(N * h_k)(\tau\Delta) := \frac{1}{T+1} \sum_{\tau'=0}^T N(\tau'\Delta) h_k((\tau - \tau')\Delta), \quad \text{for } \tau = 0, 1, \dots, T.$$

We define the Fourier transform of f as $\hat{f}(\xi) := \frac{1}{T+1} \sum_{\tau=0}^T f(\tau\Delta) e^{-2\pi i \xi(\tau\Delta)}$ for $\xi \in \mathbb{R}$. $\hat{f}(\xi)$ is a periodic function with period $1/\Delta$. Throughout this paper, $\widehat{(\cdot)}$ denotes the Fourier transform. Additionally, we assume $\mathbb{E}\beta_k =$

$\mathbb{E}R_k(t) = 0$, for all k and t , motivated by the centralization

$$Y_k(\omega; t) - \mathbb{E}Y_k(t) = \{\beta_k(\omega) - \mathbb{E}\beta_k\} (N * h_k)(t - t_{0,k}) + \{R_k(\omega; t) - \mathbb{E}R_k(t)\}.$$

Using $\{\beta_k(\omega) - \mathbb{E}\beta_k\}$ and $\{R_k(\omega; t) - \mathbb{E}R_k(t)\}$ as $\beta_k(\omega)$ and $R_k(\omega; t)$, respectively, we model the demeaned signals $Y_k(\omega; t) - \mathbb{E}Y_k(t)$. Because of the invariance $\text{corr}(\beta_k, \beta_l) = \text{corr}((\beta_k - \mathbb{E}\beta_k), (\beta_l - \mathbb{E}\beta_l))$, the assumption $\mathbb{E}\beta_k = 0$ does not prevent the detection of ptFC $|\text{corr}(\beta_k, \beta_l)|$.

Our proposed ptFC depends on neither latency $t_{0,k}$ nor reference terms $R_k(\omega; t)$. We note that, applying periodic extension (2.3.1) to $(N * h_k)$, one may verify that the two stochastic processes $(N * h_k)(t - t_{0,k} - U(\omega))$ and $(N * h_k)(t - U(\omega))$, where the random variable $U : \Omega \rightarrow \mathcal{T}$ is uniformly distributed on $\mathcal{T}$, are identically distributed. Then $(N * h_k)(t - U(\omega))$, and hence the distribution of $(N * h_k)(t - t_{0,k} - U(\omega))$, does not depend on $t_{0,k}$. Therefore, we investigate the time-shifted signals $Y_k(\omega; t - U(\omega)) = \beta_k(\omega)(N * h_k)(t - t_{0,k} - U(\omega)) + R_k(\omega; t - U(\omega))$, for $k = 1, \dots, K$. Based on the discussion above, the distributions of $Y_k(\omega; t - U(\omega))$ do not depend on $t_{0,k}$. Then, we consider the autocovariance $\mathbb{E}\{Y_k(t - U)Y_l(t + s - U)\}$ when estimating $|\text{corr}(\beta_k, \beta_l)|$, where $s = \underline{s}\Delta$ with $\underline{s} \in \mathbb{Z}$. Additionally, $|\text{corr}(\beta_k, \beta_l)|$ depends only on task-evoked terms $\{Y_k(\omega; t - U(\omega)) - R_k(\omega; t - U(\omega))\}$ (see (2.2.2)). Hence, we investigate the autocovariance difference $\mathbb{E}\{Y_k(t - U)Y_l(t + s - U)\} - \mathbb{E}\{R_k(t - U)R_l(t + s - U)\}$. If $U(\omega)$, $\{\beta_k(\omega)\}_{k=1}^K$, and $\{R_k(\omega; t)|t \in \mathcal{T}\}_{k=1}^K$ are independent, Web Theorem 2.B.1 implies that this difference depends only on s , rather than t . Therefore, we denote it as follows.

$$\mathcal{A}_{kl}(s) = \mathbb{E}\{Y_k(t - U)Y_l(t + s - U)\} - \mathbb{E}\{R_k(t - U)R_l(t + s - U)\}, \quad k, l = 1, \dots, K. \quad (2.3.2)$$

We propose the following for $\xi \in \mathbb{R}$

$$\mathcal{C}_{kl}(\xi) := |\widehat{\mathcal{A}_{kl}}(\xi)| / \sqrt{|\widehat{\mathcal{A}_{kk}}(\xi)\widehat{\mathcal{A}_{ll}}(\xi)|}$$

by incorporating a normalization multiplier in (2.3.2). Web Theorem 2.B.1 implies the following representation of ptFC.

$$\mathcal{C}_{kl}(\xi) = |\text{corr}(\beta_k, \beta_l)|, \quad \text{for } \xi \in \mathbb{R}. \quad (2.3.3)$$

The representation (2.3.3) motivates the ptFCE algorithm. Instead of estimating $|\text{corr}(\beta_k, \beta_l)|$, we propose estimating $\mathcal{C}_{kl}(\xi)$. The main steps of our proposed ptFCE algorithm provide estimation of factors $\mathcal{A}_{kl}(s)$ of $\mathcal{C}_{kl}(\xi)$, for $k, l \in \{1, \dots, K\}$.

BOLD signals $Y_k(\omega; t)$ are observed in experiments, and $U(\omega)$ is an auxiliary variable artificially generated in our estimation procedure (step 1 of Algorithm 3). Hence, the first term of (2.3.2) can be estimated using the

method of moments. However, the reference signals $R_k(\omega; t)$ implicitly contained in $Y_k(\omega; t) = P_k(\omega; t) + R_k(\omega; t)$ are not observable. Section 2.3.1 provides the estimation of $R_k(\omega; t)$.

2.3.1 The Estimation of Reference Signals

In this subsection, we derive an approximation $\tilde{R}_k(\omega; t) \approx R_k(\omega; t)$ from observed $Y_k(\omega; t)$. We start by introducing the following concept: a stochastic process $\mathbf{G}(\omega; t) = (G_1(\omega; t), \dots, G_K(\omega; t))^T$ is called *weakly stationary with mean zero* (WSMZ) if $\mathbb{E}G_k(t) = 0$, for all $t \in \mathcal{T}$, and $\mathbb{E}\{G_k(t)G_l(t+s)\}$ depends only on s , rather than t , for all $k, l \in \{1, \dots, K\}$.

First, we note that $C_k := \mathbb{E}\{(N * h_k)(t - t_{0,k} - U)\}$ is a constant depending on neither t nor $t_{0,k}$. Define stochastic processes $J_k(\omega; t) = \beta_k(\omega)\{(N * h_k)(t - t_{0,k} - U(\omega)) - C_k\}$, for $k = 1, \dots, K$. For each fixed k , one can verify that scalar-valued stochastic process $J_k(\omega; t)$ is WSMZ conditioning on β_k , then $\mathbb{E}\{J_k(t)J_k(t+s)|\beta_k\}$ depends only on s . The following theorem gives the foundation for the estimation of reference terms $R_k(\omega; t)$.

Theorem 2.3.1. *For each $k \in \{1, \dots, K\}$, suppose that BOLD signals $Y_k(\omega; t)$ are of form (2.2.2), the scalar-valued stochastic process $\{R_k(\omega; t)\}_{t \in \mathcal{T}}$ is WSMZ, and the random variable $U : \Omega \rightarrow \mathcal{T}$ is uniformly distributed. Additionally, for each k , $\beta_k(\omega)$, $\{R_k(\omega; t)\}_{t \in \mathcal{T}}$, and $U(\omega)$ are independent, and there exist $t^* \in \mathcal{T}$ and $\beta^* \in \mathbb{R}$ so that $\frac{\mathbb{E}\{J_k(t)J_k(t-t^*)|\beta_k=\beta^*\}}{\mathbb{E}\{J_k(t)^2|\beta_k=\beta^*\}} \neq \frac{\mathbb{E}\{R_k(t-U)R_k(t-t^*-U)\}}{\mathbb{E}\{R_k(t-U)^2\}}$. If nonsingular matrix $\mathbf{A} = (a_{ij})_{1 \leq i, j \leq 2}$ and stochastic process $\{\mathbf{s}(\omega; t) = (s_1(\omega; t), s_2(\omega; t))^T\}_{t \in \mathcal{T}}$ satisfy (i) given $\beta_k(\omega) = \beta^*$, $\mathbf{s}(\omega; t)$ is WSMZ, (ii) $s_1(\omega; t_1)$ and $s_2(\omega; t_2)$ are uncorrelated for $t_1, t_2 \in \mathcal{T}$, (iii) $\frac{\mathbb{E}\{s_1(t)s_1(t-t^*)|\beta_k=\beta^*\}}{\mathbb{E}\{s_1(t)^2|\beta_k=\beta^*\}} \neq \frac{\mathbb{E}\{s_2(t)s_2(t-t^*)|\beta_k=\beta^*\}}{\mathbb{E}\{s_2(t)^2|\beta_k=\beta^*\}}$, and $\mathbf{A}(s_1(\omega; t), s_2(\omega; t))^T = (Y_k(\omega; t - U(\omega)) - \beta^*C_k, (N * h_k)(t - t_{0,k} - U(\omega)) - C_k)^T$ for $t \in \mathcal{T}$ then there exist a non-singular diagonal matrix $\mathbf{\Lambda}$ and a permutation matrix $\mathbf{P}$, such that*

$$\begin{pmatrix} 1 & 1 \\ \frac{1}{\beta^*} & 0 \end{pmatrix} = \mathbf{A}\mathbf{\Lambda}^{-1}\mathbf{P}^{-1} \text{ and } \begin{pmatrix} J_k(\omega; t) \\ R_k(\omega; t - U(\omega)) \end{pmatrix} = \mathbf{P}\mathbf{\Lambda}\mathbf{s}(\omega; t), \text{ for all } t \in \mathcal{T}. \quad (2.3.4)$$

Theorem 2.3.1 herein is a straightforward result of the Theorem 2 in Tong et al. (1991). For each fixed k and ω with setting $\beta^* = \beta_k(\omega)$, the pair $(\mathbf{A}, \mathbf{s}(\omega; t))$ in Theorem 2.3.1 is derived by the AMUSE algorithm with 2D signal

$$\left\{ (Y_k(\omega; t - U(\omega)) - \beta^*C_k, (N * h_k)(t - t_{0,k} - U(\omega)) - C_k)^T \right\}_{t \in \mathcal{T}}$$

as its input. However, coefficient $\beta^* = \beta_k(\omega)$ is unknown. Since $\mathbb{E}\beta_k = 0$, we ignore $\beta_k^*C_k$ and apply the following vector-valued signal as input of the AMUSE algorithm.

$$\left\{ (Y_k(\omega; t - U(\omega)), (N * h_k)(t - t_{0,k} - U(\omega)) - C_k)^T \mid t \in \mathcal{T} \right\}. \quad (2.3.5)$$

The choice of latency $t_{0,k}$ and HRFs h_k in (2.3.5) is of importance. In our implementation of (2.3.5) in HCP

analysis, we choose $t_{0,k} = 0$, for all k , based on the following considerations: (i) latency times are usually much smaller than their corresponding experimental time units Δ (e.g., see [Zhang et al. \(2013\)](#) (p 138)); (ii) latency $t_{0,k}$ can be incorporated into the corresponding HRF h_k as its parameter, i.e., we view $h_k(\cdot - t_{0,k})$ as an HRF, then the choice of $t_{0,k}$ is equivalent to the choice of corresponding HRF; (iii) The HCP data are from a block design; if latency $t_{0,k}$ are much shorter than the length of task blocks, $t_{0,k}$ are negligible. Additionally, we choose h_k in (2.3.5) to be the canonical HRF (the R function `canonicalHRF` with default parameters in package `neuRosim`), for all k , in our analysis of the HCP data, since in block-design experiments, the influence of biased HRFs on the ptFCE algorithm is moderate as we illustrate using simulations in Section 2.4; this holds even if the support of biased HRFs is longer than task blocks. Finally, $t_{0,k}$ and h_k can be estimated using data-adaptive approaches, e.g., the spline-based method by [Zhang et al. \(2013\)](#). In general, given an algorithm for estimating $t_{0,k}$ and h_k , one may simply use the estimated $t_{0,k}^{est}$ and h_k^{est} as input in (2.3.5). Since the distribution of $U(\omega)$ is known, constants C_k can be estimated. To approximately recover $R_k(\omega; t)$ from $(\mathbf{A}, \mathbf{s}(\omega; t))$, we show the following result.

Theorem 2.3.2. *For each $k \in \{1, \dots, K\}$ and $\omega \in \Omega$, suppose the pair $(\mathbf{A}, \mathbf{s}(\omega; t))$ satisfies (2.3.4). Then there exists $i' \in \{1, 2\}$ such that $a_{1i'} s_{i'}(\omega; t) = J_k(\omega; t)$.*

The proof of Theorem 2.3.2 is in Appendix B. The index i' in Theorem 2.3.2 can be computed as

$$i' = \underset{i}{\operatorname{argmax}} \left\{ \operatorname{corr}(s_i(\omega; t), (N * h_k)(t - t_{0,k} - U(\omega_j)) - C_k), \text{ across } t | i = 1, 2 \right\}.$$

The AMUSE algorithm and Theorem 2.3.2 recover $J_k(\omega; t) = \beta_k(\omega)(N * h_k)(t - t_{0,k} - U(\omega)) - \beta_k(\omega)C_k$. Again, we ignore $\beta_k(\omega)C_k$ as $\mathbb{E}\beta_k = 0$, that is, $J_k(\omega; t) \approx \beta_k(\omega)(N * h_k)(\omega; t - t_{0,k} - U(\omega))$. Then we estimate $R_k(\omega; t - U(\omega))$ by $\tilde{R}_k(\omega; t - U(\omega)) := Y_k(\omega; t - U(\omega)) - J_k(\omega; t) \approx R_k(\omega; t - U(\omega))$, for all k and ω . In applications, $U(\omega)$ are artificially generated and known. Then we have the approximation $\tilde{R}_k(\omega; t) \approx R_k(\omega; t)$. We conclude the derivation of $\tilde{R}_k(\omega; t)$ from $Y_k(\omega; t)$ as follows: Step 1, each observed signal $Y_k(\omega; t)$ provides input (2.3.5) for the AMUSE algorithm; Step 2, the AMUSE algorithm computes $(\mathbf{A}, \mathbf{s}(\omega; t))$; Step 3, Theorem 2.3.2 indicates $J_k(\omega; t) = a_{1i'} s_{i'}(\omega; t)$; Step 4, compute $\tilde{R}_k(\omega; t)$ using $J_k(\omega; t)$ and the artificially generated $U(\omega)$.

2.3.2 The Estimator for $\mathcal{C}_{kl}(\xi)$ and The ptFCE Algorithm

With observed signals $\{Y_{k'}(\omega; t) | k' = k, l\}_{\omega=1}^n$ and the signals $\{\tilde{R}_{k'}(\omega; t) | k' = k, l\}_{\omega=1}^n$ derived from $\{Y_{k'}(\omega; t) | k' = k, l\}_{\omega=1}^n$, we propose the following estimator for $\mathcal{C}_{kl}(\xi) = |\text{corr}(\beta_k, \beta_l)|$.

$$\begin{aligned} \mathcal{C}_{kl}^{est,n}(\xi) &:= \left| \widehat{\mathcal{A}_{kl}^{est,n}}(\xi) \right| / \sqrt{\left| \widehat{\mathcal{A}_{kk}^{est,n}}(\xi) \widehat{\mathcal{A}_{ll}^{est,n}}(\xi) \right|}, \quad \text{for all } \xi \in \mathbb{R}, \text{ where} \\ \mathcal{A}_{kl}^{est,n}(s) &:= \underline{Y}(s) - \underline{\tilde{R}}(s), \\ \underline{Y}(s) &:= \frac{1}{T+1} \sum_{t \in \{\tau\Delta\}_{\tau=0}^T} \left[\frac{1}{n} \sum_{\omega=1}^n \{Y_k(\omega; t - U(\omega)) Y_l(\omega; t + s - U(\omega))\} \right], \\ \underline{\tilde{R}}(s) &:= \frac{1}{T+1} \sum_{t \in \{\tau\Delta\}_{\tau=0}^T} \left[\frac{1}{n} \sum_{\omega=1}^n \{\tilde{R}_k(\omega; t - U(\omega)) \tilde{R}_l(\omega; t + s - U(\omega))\} \right]. \end{aligned} \quad (2.3.6)$$

The periodicity property of the Fourier transform implies that $\mathcal{C}_{kl}^{est,n}(\xi)$ is a periodic function of ξ with period $1/\Delta$. Web Lemma 2.B.1 implies that (2.3.6) is well-defined for all $\xi \in [0, 1/\Delta]$ except for at most finitely many points ξ such that $\widehat{\mathcal{A}_{kk}^{est,n}}(\xi) \widehat{\mathcal{A}_{ll}^{est,n}}(\xi) = 0$.

$\mathcal{A}_{kl}^{est,n}(s)$ has a method of moments form except R_k in (2.3.2) is replaced with the estimated $\tilde{R}_k$. This has an effect on the consistency of $\mathcal{C}_{kl}^{est,n}(\xi)$ as an estimator for $\mathcal{C}_{kl}(\xi)$ and results in the asymptotic bias $|\lim_{n \rightarrow \infty} \mathcal{C}_{kl}^{est,n}(\xi) - \mathcal{C}_{kl}(\xi)|$. The bias can be small for properly chosen ξ . Specifically, we derive a formula of the bias as a function of ξ and choose ξ that minimizes the bias. Since $\tilde{R}_k(\omega; t)$ approximates $R_k(\omega; t)$, we model the difference $W_k(\omega; t) := R_k(\omega; t) - \tilde{R}_k(\omega; t)$ as random noise and assume that $\mathbf{W}(\omega; t) := (W_1(\omega; t), \dots, W_K(\omega; t))^T$ satisfies (i) $\mathbf{W}(\omega; t_1)$ is independent of $\mathbf{W}(\omega; t_2)$ if $t_1 \neq t_2$, (ii) $\mathbf{W}(\omega; t)$ is WSMZ, and (iii) $\Sigma_{kl} := \mathbb{E}[W_k(t)W_l(t)]$ for $t \in \mathcal{T}$. Because of $\Sigma_{kl} \leq \sqrt{\Sigma_{kk}\Sigma_{ll}}$ and $\mathbb{E}\beta_k = \mathbb{E}\beta_l = 0$, under the assumptions on $\mathbf{W}(\omega; t)$, Web Theorem 2.B.2 implies

$$\lim_{n \rightarrow \infty} \mathcal{C}_{kl}^{est,n}(\xi) \approx \frac{\left| \mathbb{E}(\beta_k \beta_l) \frac{\widehat{h}_k(\xi) \widehat{h}_l(\xi)}{|\widehat{h}_k(\xi) \widehat{h}_l(\xi)|} e^{2\pi i \xi(t_{0,k} - t_{0,l})} \right|}{\sqrt{\mathbb{E}(\beta_k^2) \mathbb{E}(\beta_l^2)}} = \frac{|\mathbb{E}(\beta_k \beta_l)|}{\sqrt{\mathbb{E}(\beta_k^2) \mathbb{E}(\beta_l^2)}} = |\text{corr}(\beta_k, \beta_l)|$$

almost surely at frequencies ξ such that

$$\Sigma_{kk} / [(T+1)|\widehat{N}(\xi) \widehat{h}_k(\xi)|^2] \approx 0, \quad \Sigma_{ll} / [(T+1)|\widehat{N}(\xi) \widehat{h}_l(\xi)|^2] \approx 0. \quad (2.3.7)$$

The ξ resulting in sufficiently large $|\widehat{h}_k(\xi)|$ and $|\widehat{h}_l(\xi)|$ satisfying (2.3.7) reduces the estimation bias. From the signal processing perspective, large $|\widehat{h}_k(\xi)|$ and $|\widehat{h}_l(\xi)|$ filter out the random noise in $W_k(\omega; t)$, and HRFs act as band-pass filters (typically 0 – 0.15 Hz, Aguirre et al. (1997)). Hence, we are interested in $\xi \in (0, 0.15)$. Therefore, motivated by (2.3.7), $|\text{corr}(\beta_k, \beta_l)|$ is approximated by the median of $\mathcal{C}_{kl}^{est,n}(\xi)$ across $\xi \in (0, 0.15)$, since the median is more stable than mean. One typical curve of $\{\mathcal{C}_{kl}^{est,n}(\xi) | \xi \in \mathbb{R}\}$ is in Figure 2.1. The estimator $\mathcal{C}_{kl}^{est,n}(\xi)$ in (2.3.6) and the choice of ξ complete the estimation of $\mathcal{C}_{kl}(\xi) = |\text{corr}(\beta_k, \beta_l)|$. The estimation procedure is concluded in the

ptFCE algorithm (Algorithm 3).

Algorithm 3 ptFCE Algorithm

Input: (i) Task-fMRI BOLD signals $\{(Y_k(\omega; \tau\Delta), Y_l(\omega; \tau\Delta))\}_{\tau=0}^T$ for sampled subjects $\omega \in \{1, \dots, n\}$; (ii) repetition time Δ ; (iii) the stimulus signal $\{N(\tau\Delta)\}_{\tau=0}^T$; (iv) latency $t_{0,k'} = \tau_{0,k'}\Delta$ with some $\tau_{0,k'} \in \mathbb{Z}$ for $k' \in \{k, l\}$, and their default values are $t_{0,k'} = 0$ for $k' \in \{k, l\}$; (v) HRF $h_{k'}$ for $k' \in \{k, l\}$. The default HRFs are $h_k = h_l =$ the R function `canonicalHRF` with its default parameters.

Output: An estimation of the ptFC $|corr(\beta_k, \beta_l)|$ between the k^{th} and l^{th} nodes.

- 1: Generate i.i.d. $\{U(\omega)\}_{\omega=1}^n$ from the uniform distribution on $\mathcal{T} = \{\tau\Delta\}_{\tau=0}^T$.
 - 2: For each $\omega \in \{1, \dots, n\}$ and $k' \in \{k, l\}$, apply the AMUSE algorithm to input (2.3.5) and obtain estimated reference signals $\{\tilde{R}_k(\omega; \tau\Delta)\}_{\tau=0}^T$.
 - 3: Compute the estimator $\mathcal{C}_{kl}^{est,n}(\xi)$ in (2.3.6) and the median of $\{\mathcal{C}_{kl}^{est,n}(\xi) | \xi \in (0, 0.15)\}$. The median is the output of this algorithm.
-

2.4 Simulations

When comparing methods for FC estimation, direct comparisons of methods are difficult as the estimates are defined on different scales (Cisler et al. (2014)). In the context of FC analysis, we may compare the methods by measuring their accuracy in correctly identifying the relative values of FC among pairs of nodes estimated by each approach, i.e., the performance of the algorithm in terms of detecting underlying “weak vs. strong” patterns instead of the estimated values. Specifically, suppose we investigate two pairs of nodes. The connectivity between one pair is weak while the connectivity between another pair is strong; we are interested in whether the estimated value for the weak connectivity is smaller than that of the strong connectivity. Therefore, in this section, we only compare ptFCE with other approaches in detecting underlying “weak vs. strong” patterns. Simulation studies showing the performance of ptFCE in terms of estimation bias and variance are presented in Appendix D. For the comparative analysis, we use synthetic data generated by two different mechanisms designed to mimic the properties of HCP motor data. We design the simulations by using non-canonical HRFs in generating synthetic data to illustrate the robustness of our algorithm that uses the canonical HRF as input. We show that the algorithm can still detect underlying task-evoked FC patterns.

We compare the ptFCE algorithm with the following methods: beta-series regression, *naive Pearson correlation*, *task Pearson correlation*, coherence analysis using two data-generating mechanisms such that the first is based on model (2.2.2) for defining ptFCs, while the second is motivated by Pearson correlations. We use the second data generating mechanism to illustrate the comparable performance of our proposed approach with others even when the data generating mechanism is not compatible with the ptFC framework.

Mechanism 1: Step 1, generate coefficients $(\beta_1(\omega), \beta_2(\omega), \beta_3(\omega))^T \sim N_3(\mathbf{0}, (\sigma_{ij})_{1 \leq i, j \leq 3})$ for $\omega = 1, \dots, n$. Assume $\rho_{ij} := |\sigma_{ij} / \sqrt{\sigma_{ii}\sigma_{jj}}|$, for $(i, j) \in \{(1, 2), (2, 3)\}$ are the true underlying ptFCs between nodes 1 and 2 and between nodes 2 and 3, respectively. **Step 2**, generate reference signals $\{R_{k'}(\omega; \tau\Delta) = Q_{k'}(\omega; \tau\Delta) + \epsilon_{k'}(\omega; \tau\Delta)\}_{\tau=0}^T$, for all ω and $k' \in \{1, 2, 3\}$, where the $3nT$ noise values $\{\epsilon_{k'}(\omega; \tau\Delta) | k' = 1, 2, 3; \omega = 1, \dots, n; \tau = 0, \dots, T\} \sim_{iid}$

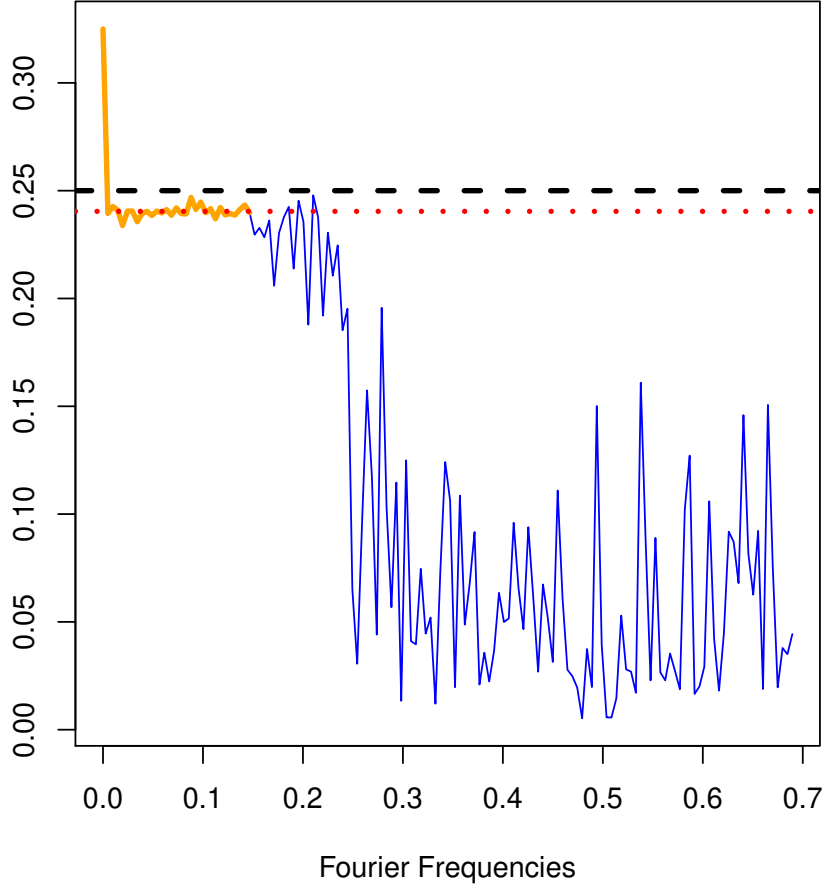


Figure 2.1: To generate this figure, we first generate synthetic BOLD signals using the Mechanism 0 in Appendix E with sample size $n = 308$ and underlying $\rho = 0.25$ (presented by the black dashed line). Then we apply our proposed ptFCE algorithm to the synthetic BOLD signals generated by Mechanism 0. The solid blue curve presents the estimator $\mathcal{C}_{kl}^{est,n}(\xi)$ as a function of Fourier frequencies $\xi \in (0, \frac{1}{2\Delta})$, and the solid orange curve presents $\{\mathcal{C}_{kl}^{est,n}(\xi) | \xi \in (0, 0.15)\}$ for taking the desired median. The unreasonably large value part of this orange curve motivates us to implement the median instead of the mean of $\{\mathcal{C}_{kl}^{est,n}(\xi) | \xi \in (0, 0.15)\}$. The dotted red line presents the median of $\mathcal{C}_{kl}^{est,n}(\xi)$ across $\xi \in (0, 0.15)$, i.e., the output of the ptFCE algorithm for the synthetic BOLD signals.

$N(0, V)$. **Step 3**, compute synthetic BOLD signals $\{Y_{k'}(\omega; \tau\Delta) | k' = 1, 2, 3\}_{\tau=0}^T$, for $\omega = 1, \dots, n$, by $Y_{k'}(\omega; \tau\Delta) = 9000 + \beta_{k'}(\omega) \times (N * h_{k'}) (\tau\Delta) + R_{k'}(\omega; \tau\Delta)$. $\mathbf{Y}(\omega; t) := \{Y_{k'}(\omega; \tau\Delta)\}_{k'=1}^3$.

Mechanism 2: Step 1, for each $\omega \in \{1, \dots, n\}$, generate independent normal vectors

$$\{\boldsymbol{\varepsilon}(\omega; \tau\Delta) = (\varepsilon_1(\omega; \tau\Delta), \varepsilon_2(\omega; \tau\Delta), \varepsilon_3(\omega; \tau\Delta))^T\}_{\tau=0}^T,$$

where $\boldsymbol{\varepsilon}(\omega; \tau\Delta) \sim N_3(\mathbf{0}, \boldsymbol{\Sigma}_{task})$ if $N(\tau\Delta) = 1$, and $\boldsymbol{\varepsilon}(\omega; \tau\Delta) \sim N_3(\mathbf{0}, \boldsymbol{\Sigma}_{resting})$ if $N(\tau\Delta) = 0$; matrices $\boldsymbol{\Sigma}_{task}$ and $\boldsymbol{\Sigma}_{resting}$ share the same diagonals, the off-diagonal elements of $\boldsymbol{\Sigma}_{resting}$ are 0, and $\{\varrho_{ij}\}_{i,j=1}^3$ denote the correlations deduced from covariance matrix $\boldsymbol{\Sigma}_{task}$. **Step 2**, compute $Y_{k'}(\omega; \tau\Delta) = N * h_{k'}(\tau\Delta) + \sum_{j=1}^4 N_j * h_{k',j}(\tau\Delta) + \varepsilon_{k'}(\tau\Delta)$, for $k' \in \{1, 2, 3\}$, where the tasks $\{N_j(t)\}_{j=1}^4$ are nuisance tasks, and $\{h_{k',j}(t)\}_{j=1}^4$ are the corresponding HRFs.

Step 3, repeat Steps 1 and 2 for all $\omega = 1, \dots, n$.

Details of Mechanisms 1 and 2 are in Appendix E. In Mechanism 2, the zero off-diagonal elements of $\mathbf{\Sigma}_{resting}$ indicate no task-evoked connectivity when the task of interest is absent. When $N(\tau\Delta) = 1$, the correlation ϱ_{ij} in $\mathbf{\Sigma}_{task}$ measures the connectivity between nodes i and j evoked by task $N(\tau\Delta) = 1$. We are interested in estimating the connectivity between nodes 1, 2 and nodes 2, 3. With the synthetic signals $\{\mathbf{Y}(\omega; t)\}_{\omega=1}^n$ generated by Mechanisms 1 or 2, the existing methods are implemented as follows.

Naive Pearson correlation: For each subject ω and underlying ρ_{ij} or ϱ_{ij} , compute the Pearson correlation between $\{Y_i(\omega; \tau\Delta)\}_{\tau=0}^T$ and $\{Y_j(\omega; \tau\Delta)\}_{\tau=0}^T$ across all τ and denote the absolute value of this correlation by $\hat{\rho}_{ij,\omega}^{naiveCorr}$. Let $\hat{\rho}_{ij,mean}^{naiveCorr}$ and $\hat{\rho}_{ij,median}^{naiveCorr}$ denote the mean and median, respectively, of $\{\hat{\rho}_{ij,\omega}^{naiveCorr}\}_{\omega=1}^n$ across all ω .

Task Pearson correlation: For each ω and underlying ρ_{ij} or ϱ_{ij} , compute the Pearson correlation between $\{Y_i(\omega; \tau\Delta)|N(\tau\Delta) = 1\}$ and $\{Y_j(\omega; \tau\Delta)|N(\tau\Delta) = 1\}$ across τ such that $N(\tau\Delta) = 1$ and denote its absolute value by $\hat{\rho}_{ij,\omega}^{taskCorr}$. Compute the mean and median of $\{\hat{\rho}_{ij,\omega}^{taskCorr}\}_{\omega=1}^n$ across ω and denote them by $\hat{\rho}_{ij,mean}^{taskCorr}$ and $\hat{\rho}_{ij,median}^{taskCorr}$, respectively.

Details of implementing the beta-series regression and coherence analysis to obtain estimates $(\hat{\rho}_{ij,mean}^{betaS}, \hat{\rho}_{ij,median}^{betaS})$ and $(\hat{\rho}_{ij,mean}^{Coh}, \hat{\rho}_{ij,median}^{Coh})$ are presented in Appendix G.

Let $(\rho_{12}, \rho_{23}) = (\varrho_{12}, \varrho_{23}) = (0.4, 0.6)$, indicating weak (0.4) or strong (0.6) connectivity between the two node pairs. We use sample sizes $n = 50$ and 308 to illustrate the performance in different sample sizes, where 308 is the size of the HCP data set. PtFCs estimated by the ptFCE algorithm are denoted by $(\hat{\rho}_{12}, \hat{\rho}_{23})$. For each simulated data set, applying all FC estimation methods, we obtain estimates $\{(\hat{\rho}_{ij}, \hat{\rho}_{ij,mean}^{betaS}, \dots, \hat{\rho}_{ij,median}^{Coh})|i < j\}_{i,j=1}^3$. For each method, we evaluate whether it can identify the “weak vs. strong” pattern. For example, our proposed ptFCE algorithm is effective if $\hat{\rho}_{12} < \hat{\rho}_{23}$. We repeat this procedure 500 times. The rates of correct identification of the connectivity patterns across 500 simulations are presented in Table 2.1. When synthetic BOLD signals are from the mechanism compatible with the definition of ptFCs - Mechanism 1, our ptFCE algorithm performs overwhelmingly better than other methods. If the synthetic signals are from the mechanism motivated by Pearson correlations - Mechanism 2, our proposed algorithm still provides a good identification rate and is better than the coherence analysis. Importantly, it is unlikely that the true data generating process in task-fMRI follows the second data generating process. Mechanism 2 implies the effect of task disappears as soon as the task is absent. However, it is known that there is a delay in brain response associated with time needed by brain vasculature to respond to the decrease in oxygen (Lindquist et al. (2008)). Even in this unrealistic scenario, our proposed method is comparable with others. Naive Pearson correlation and coherence analysis are designed to measure FC evoked by both the task of interest and nuisance tasks. In contrast, beta-series regression, task Pearson correlations, and ptFCE are developed for quantifying FC evoked by specific tasks of interest. Hence, only beta-series regression, task Pearson correlations, and ptFCE algorithms are useful when estimating task-evoked FC. Additionally, task Pearson correlation model assumes that the effect of task of interest disappears immediately when the task is

Table 2.1: Identification rates $\hat{p} := n'/n$ of different methods for the “weak vs strong” pattern, where n' is the number of correct identifications among all $n = 500$ synthetic samples. The correct identification rates in the table are presented in 95%-confidence interval form $\hat{p} \pm z_{0.975} \cdot \sqrt{\hat{p}(1 - \hat{p})/n}$, where $z_{0.975} \approx 1.96$ is the 0.975-quantile of the standard normal distribution $N(0, 1)$.

Methods	Rates ($n = 50$)		Rates ($n = 308$)	
	Mech. 1	Mech. 2	Mech. 1	Mech. 2
ptFCE	82.8(± 10.5)%	67.4(± 13.0)%	99.8(± 0.5)%	88.4(± 2.9)%
SLR	66.4(± 13.1)%	56.8(± 13.7)%	84.8(± 3.6)%	63.4(± 5.4)%
naive Pearson (mean)	58.2(± 13.7)%	93.8(± 6.7)%	69.0(± 5.2)%	100(± 0)%
naive Pearson (median)	55.6(± 13.8)%	88.0(± 9.0)%	62.2(± 5.4)%	100(± 0)%
task Pearson (mean)	47.6(± 13.8)%	100(± 0)%	48.6(± 5.6)%	100(± 0)%
task Pearson (median)	47.8(± 13.8)%	100(± 0)%	50.6(± 5.6)%	100(± 0)%
beta-series (mean)	49.4(± 13.9)%	96.6(± 5.0)%	48.8(± 5.6)%	100(± 0)%
beta-series (median)	48.2(± 13.9)%	97.8(± 4.1)%	47.6(± 5.6)%	100(± 0)%
coherence (mean)	51.0(± 13.9)%	55.8(± 13.8)%	60.8(± 5.5)%	72.0(± 5.0)%
coherence (median)	51.6(± 13.6)%	57.6(± 13.7)%	58.8(± 5.5)%	68.0(± 5.2)%

absent. Hence, because of the delay in brain response, task Pearson correlations may be biased. Lastly, Table 2.1 shows that ptFCE performs better than the beta-series regression for data from Mechanism 1. In addition, we compare the SLR approach introduced in the beginning of Section 2.3 with the ptFCE algorithm. When the synthetic data are generated from Mechanism 1, although SLR performs overwhelmingly better than the existing methods, it is inferior to the ptFCE algorithm. When the synthetic data are generated from Mechanism 2, which violates the assumptions in model (2.2.2), SLR is hardly better than a random guess.

PtFCE is computationally efficient and scalable. On a PC with 2.4 GHz 8-Core Intel Core i9 processor, the approximate computational time of ptFCE is 3.5 seconds for 50 subjects and 11.5 seconds for 1000 subjects.

2.5 Analysis of HCP Motor Data

In this section, we present estimation of FC in a task-evoked functional MRI study using data from HCP. For comparison, we apply the ptFCE algorithm and existing methods to measure FC in the database of 308 subjects from HCP performing motor tasks. A detailed description of the HCP data is provided in Section 2.1. We model squeezing right toes as the task of interest. Before the estimation step, we compute region-specific time courses using the AAL atlas (Tzourio-Mazoyer et al. (2002)) that consists of 120 brain regions. For each region, we extract the voxel-specific time series in that region and compute their spatial average for each time point. As a result, we obtain 120 time courses corresponding to 120 regions of interest. We select the left precentral gyrus (PreCG.L) as the seed region since it is located in the primary motor cortex and the motions of the right toes are associated with the left brain. We measure FC induced by $N(t)$ at a population level between the seed region and other regions using the following five approaches: the ptFCE algorithm, naive and task Pearson correlations, beta-series regression, and coherence analysis. In Section 2.4, the means $\hat{\rho}_{ij,mean}^{naiveCorr}$, $\hat{\rho}_{ij,mean}^{taskCorr}$, $\hat{\rho}_{ij,mean}^{betaS}$, and $\hat{\rho}_{ij,mean}^{Coh}$ tend to perform better than the corresponding medians. Hence, we show only the mean results in data analysis. Because

of numerical issues, we omit regions CB7.L, CB7.R, and CB10.L and investigate the rest of the 117 regions. The detailed procedures of applying these approaches are described in Section 2.4.

Suppose $\mathcal{X}_{r,j}$, for $r = 1, \dots, 116$ and $j = 1, \dots, 5$, are the estimated FC values between **PreCG.L** and other 116 regions computed using the referred five approaches indexed by j . Since these approaches use different scales, we standardize the estimates by $\mathcal{X}_{r,j}^{(st)} = \frac{\mathcal{X}_{r,j} - \min\{\mathcal{X}_{r',j}\}_{r'=1}^{116}}{\max\{\mathcal{X}_{r',j}\}_{r'=1}^{116} - \min\{\mathcal{X}_{r',j}\}_{r'=1}^{116}}$, for all r and j , to enable comparisons. The standardized versions of the estimates are presented in Figure 2.2. To quantify the agreement between estimation approaches implemented in this study, we use 0.5 as the threshold for standardized $\mathcal{X}_{r,j}^{(st)}$ to determine whether the node r is connected to **PreCG.L** according to method j . Specifically, if $\mathcal{X}_{r,j}^{(st)} > 0.5$, then method j estimates that region r is functionally connected to **PreCG.L**. Define $\mathbf{1}(\mathcal{X}_{r,j}^{(st)} > 0.5) =: \vartheta_{r,j}$. As a result of applying each method j , we obtain the connectivity pattern $\{\vartheta_{r,j}\}_{r=1}^{116}$. Agreement between methods j_1 and j_2 , for all $j_1, j_2 = 1, \dots, 5$, is measured by Cohen’s kappa statistic denoted as κ_{j_1, j_2} (Cohen (1960)) using classification results $\{\vartheta_{r, j_1}\}_{r=1}^{116}$ and $\{\vartheta_{r, j_2}\}_{r=1}^{116}$. The κ_{j_1, j_2} along with p-values testing the null hypothesis indicating that the extent of agreement between each pair of methods is the same as random ($\kappa_{j_1, j_2} = 0$) are presented in Table 2.2. Using $\alpha = 0.05$ and implementing multiple comparisons correction, we find that all methods have significant agreement with each other.

The postcentral gyrus (**PoCG.L**) region is identified by all five methods as the region with strongest functional connectivity with the seed region. The estimated ptFCs indicate that the neuronal activity of the left precentral gyrus and that of left postcentral gyrus corresponding to squeezing right toes are highly correlated (either negatively or positively). Magnitudes of the neuronal activity corresponding to the task of interest in these two regions tend to be linearly dependent across the entire population. In addition to modeling advantages of our proposed ptFC approach described in Section 2.2.2, we obtain differences between the estimates of ptFCE and those of competitor methods. We identify high connectivity between several regions and the seed region that are missed by competitor methods. Specifically, we obtain high task-evoked FC between the seed region and left/right thalamus (passing motor signals to the cerebral cortex), left/right paracentral lobule (motor nerve supply to the lower extremities), left superior temporal gyrus (containing the auditory cortex), and left Heschl gyrus (in the area of primary auditory cortex). These regions are related to the motor function or auditory cortex and can add to the existing results on functional connectivity between the precentral gyrus and the rest of the brain.

To further visualize ptFCs induced by the task of interest, we apply MNI space coordinates of the 117 regions from the AAL atlas, where three-dimensional coordinates are obtained from **aal2.120** dataset in R package **brainGraph**. The regions are depicted by their MNI coordinates in Figure 2.3. Grayscale shade of edges connecting each region to the seed region illustrates estimated ptFC between the corresponding region and the seed region (Figure 2.3, left panel), while the edges in the right panel of Figure 2.3 present the 30 highest ptFC values. Figure 2.3 shows that most of the large estimated ptFC values are in the left brain. This is expected, since it is known that behaviors of extremities are functionally associated with contralateral brain regions (Nieuwenhuys et al. (2014)).

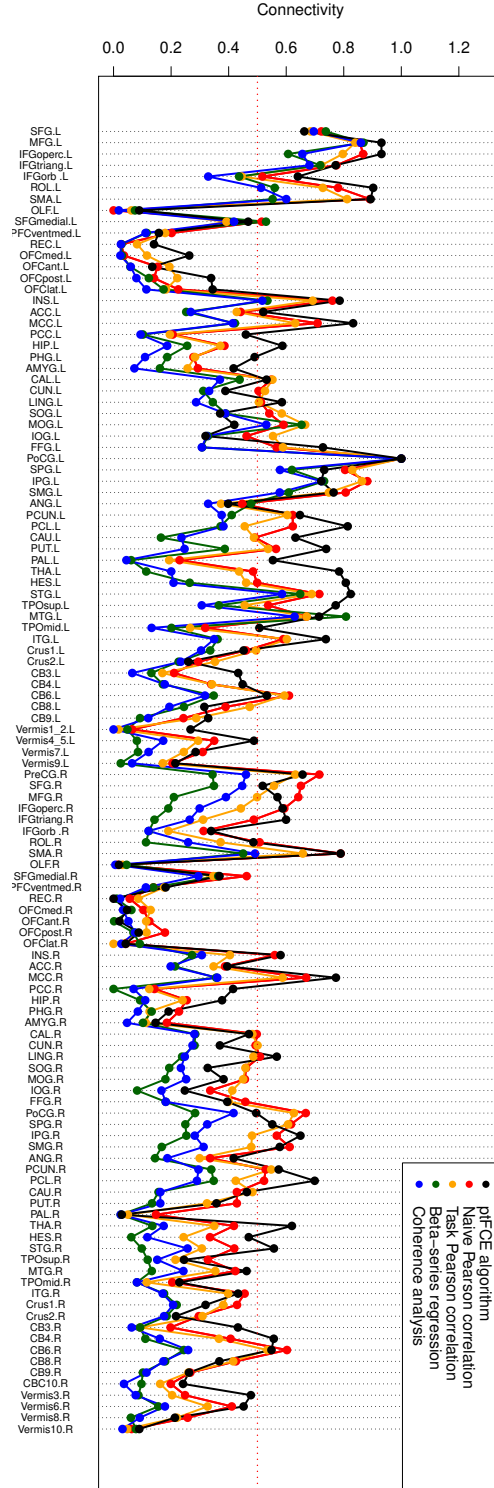


Figure 2.2: Illustration of estimation results from five FC estimation methods. The horizontal axis indicates the 116 regions compared to **PreCG.L**. The abbreviations of region names are provided in the data set **aa12.120** in the R package **brainGraph**. The vertical axis presents the standardized connectivity measurements $\chi_{r,j}^{(st)}$ between each region and the seed region **PreCG.R**. The red dotted horizontal line indicates the threshold 0.5 implemented for determining the connected/nonconnected relationships between **PreCG.L** and other regions.

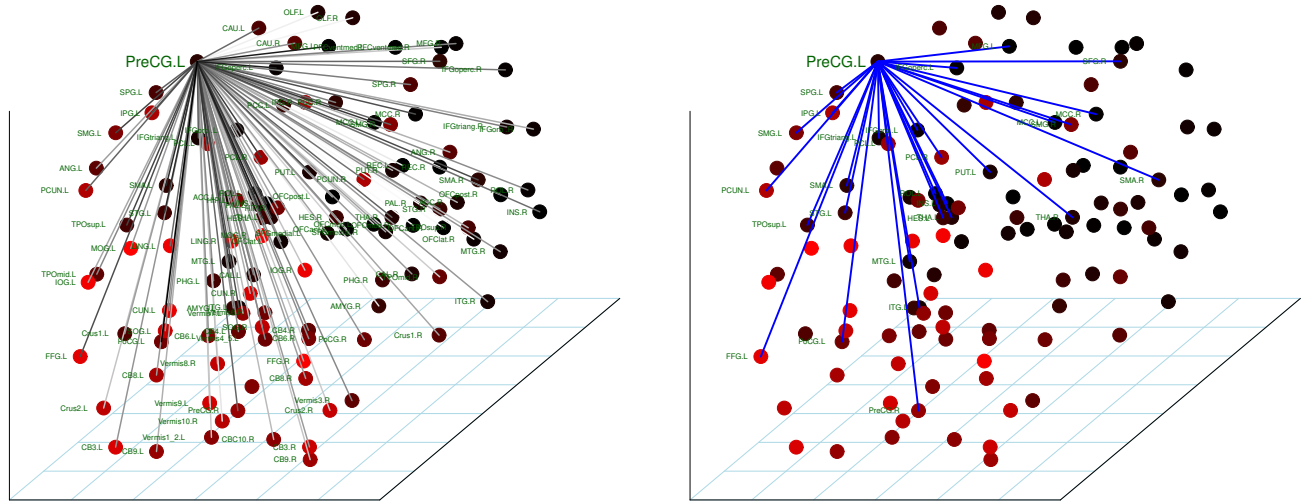


Figure 2.3: Each dot denotes a region from the AAL atlas, located using its corresponding MNI coordinates. The abbreviated region names are given next to each dot. We apply the ptFCE algorithm to estimate the ptFCs between region **PreCG.L** and each of the rest 116 regions. In the left panel, we use grayscale coloring of the edges to indicate the magnitude of ptFC between the corresponding two vertices; specifically, the larger a ptFC, the darker the line segment connecting the corresponding region pair. In the right panel, the presented blue line segments indicate the 30 largest ptFCs estimated by the ptFCE algorithm among all 116 regions.

2.6 Conclusions

In this paper, we propose a rigorous random-effects type model for task-evoked BOLD signals in fMRI studies. Based on this model, we define a measurement, ptFC, of task-evoked FC at a population level. We discuss shortcomings of several existing methods for evaluating task-evoked population-level FC that are addressed by our proposed framework. Thorough theoretical results on the properties of our proposed estimation procedure are presented. We develop the ptFCE algorithm for estimating ptFC and show that it is computationally efficient. The superior performance of our proposed ptFCE algorithm when data are generated using a block-design task

Table 2.2: The Kappa statistics between each pair of methods are computed using the R function `Kappa.test` in package `fmsb`. The p-value of a Kappa statistic is presented in the parentheses below the statistic.

Methods	ptFCE	naive Persn	task Persn	beta series	coherence
ptFCE	*	0.696 (< 0.0001)	0.503 (< 0.0001)	0.260 (0.00636)	0.277 (0.00402)
naive Persn	*	*	0.732 (< 0.0001)	0.391 (0.00022)	0.367 (0.00051)
task Persn	*	*	*	0.478 (< 0.0001)	0.497 (< 0.0001)
beta series	*	*	*	*	0.961 (< 0.0001)
coherence	*	*	*	*	*

fMRI data-generating mechanism is illustrated using simulation studies. We use ptFCE to estimate ptFCs in a motor-task data set publicly available from HCP. PtFCE results in similar FC patterns with the widely used existing methods. Additionally, ptFCE provides FC discoveries of connections between nodes not identified by existing methods. Extensions of our proposed work can include development of a subject-level task-evoked FC definition and estimation procedure, as well as considerations on the inclusion of interaction terms $P_k(\omega; t)Q_k(\omega; t)$ (see Appendix H).

Data Availability Statement

The data supporting the findings in Section 2.5 are publicly available on the HCP website:

<https://protocols.humanconnectome.org/HCP/3T/task-fMRI-protocol-details.html> (Barch et al. (2013)).

The R code for simulation studies and data analyses is available at

<https://github.com/KMengBrown/Population-level-Task-evoked-Functional-Connectivity.git>.

Appendix

2.A Relationship between the proposed BOLD signal model and some existing models

In this paper, we implement the following model for observed task-fMRI BOLD signals $Y_k(\omega; t)$, for $\omega \in \Omega$ and $k \in \{1, \dots, K\}$.

$$Y_k(\omega; t) = P_k(\omega; t) + Q_k(\omega; t) + \epsilon_k(\omega; t), \quad \text{with} \quad (2.A.1)$$

$$P_k(\omega; t) = \beta_k(\omega) \times N * h_k(t - t_{0,k}),$$

where $P_k(\omega; t)$ denotes the BOLD signal component stemming solely from the task $N(t)$ of interest, $Q_k(\omega; t)$ denotes spontaneous neuronal activity component and neuronal activity responding to nuisance tasks, and $\epsilon_k(\omega; t)$ is random error. Additionally, $t_{0,k}$ denotes the population shared latency and is usually smaller than the experimental unit, e.g., see [Zhang et al. \(2013\)](#) (p 138). The HRF shared by the whole population is $h_k(t)$. Here, we explore three existing models for BOLD signals. We show that these models are special cases of our proposed BOLD signal model (2.A.1).

Example 1: Using the notations in our paper, the BOLD signal model implemented by [Zhang et al. \(2013\)](#) can be represented as follows.

$$Y_k(\omega; t) = X_k(\omega; t)^T d(\omega) + \sum_{\gamma=1}^{\Gamma} \beta_{k,\gamma}(\omega) \times (N_\gamma * h_{k,\gamma})(t) + \epsilon_k(\omega; t), \quad (2.A.2)$$

where $N_\gamma(t)$ are task stimulus signals, $h_{k,\gamma}(t)$ are corresponding HRFs, and $X_k(\omega; t)^T d(\omega)$ characterizes the BOLD signal components stemming from other known sources, e.g., respiration and heartbeat. The BOLD signal model proposed by [Zhang et al. \(2013\)](#) includes participant-dependent latency times. These latency times are essentially modeled as zero, since the values are usually much smaller than the corresponding experimental time unit.

Suppose we are interested in the first experimental task $N_1(t)$. Define the following.

$$\beta_k(\omega) := \beta_{k,1}(\omega), \quad N(t) := N_1(t), \quad P_k(\omega; t) = \beta_k(\omega) \times (N * h_k)(t), \quad (2.A.3)$$

$$Q_k(\omega; t) := X_k(\omega; t)^T d(\omega) + \sum_{\gamma=2}^{\Gamma} \beta_{k,\gamma}(\omega) \times (N_\gamma * h_{k,\gamma})(t).$$

Then (2.A.2) is equivalent to the model (2.A.1) with latency $t_{0,k} = 0$, and there is no interaction term $P_k(\omega; t)Q_k(\omega; t)$.

Example 2: Using the notations in our paper, the BOLD signal model implemented by [Warnick et al. \(2018\)](#) is

represented as follows (see the equations (1) and (2) in [Warnick et al. \(2018\)](#)).

$$Y_k(\omega; t) = \mu_k(\omega) + \sum_{\gamma=1}^{\Gamma} \beta_{k,\gamma}(\omega) \times (N_\gamma * h_{k,\gamma})(t) + \epsilon_k(\omega; t), \quad (2.A.4)$$

where $N_\gamma(t)$ are task stimulus signals, $h_{k,\gamma}(t)$ are corresponding HRFs, $\epsilon_k(\omega; t)$ are random error, and $\mu_k(\omega)$ present the baseline. Suppose we are interested in task $N_1(t)$. We apply the transform (2.A.3) and define $Q_k(\omega; t) = \mu_k(\omega) + \sum_{\gamma=2}^{\Gamma} \beta_{k,\gamma}(\omega) \times (N_\gamma * h_{k,\gamma})(t)$. Then model (2.A.4) is equivalent to our proposed model (2.A.1) with latency $t_{0,k} = 0$, without an interaction term $P_k(\omega; t)Q_k(\omega; t)$.

Example 3: Motivated by the *independent component analysis* framework, [Joel et al. \(2011\)](#) implemented the following model for task-fMRI BOLD signals.

$$\begin{aligned} Y_k(\omega; t) = & M_{m,k}(\omega) \{ \beta_{tm}(\omega) \times N * h_k(t) + \beta_{im}(\omega) \iota_m(t) \} \\ & + M_{v,k}(\omega) \{ \beta_{tv}(\omega) \times N * h_k(t) + \beta_{iv}(\omega) \iota_v(t) \} \\ & + M_b(\omega) \beta_b(\omega) \gamma(\omega; t), \end{aligned} \quad (2.A.5)$$

where $M_{(\cdot),k}$ is the spatial mask at the k^{th} node for visual cortex ($M_{v,k}$), motor cortex ($M_{m,k}$) or the whole brain (M_b), $N(t)$ is stimulus signal corresponding to the task of interest, $h_k(t)$ is an HRF, $\iota_{(\cdot)}$ is the intrinsic activity of the corresponding cortex, $\gamma(\omega; t)$ presents random noise, and $\beta_{tm}(\omega), \beta_{im}(\omega), \beta_{tv}(\omega), \beta_{iv}(\omega), \beta_b(\omega)$ are the weights of motor task, intrinsic motor activity, visual task, intrinsic visual activity and noise, respectively.

Since the task of interest in the HCP data is a motor task, we define the following notations.

$$\begin{aligned} \beta_k(\omega) &:= M_{m,k}(\omega) \times \beta_{tm}(\omega), \\ Q_k(\omega) &:= M_{m,k}(\omega) \beta_{im}(\omega) \iota_m(t) + M_{v,k}(\omega) \{ \beta_{tv}(\omega) \times N * h_k(t) + \beta_{iv}(\omega) \iota_v(t) \}, \\ \epsilon_k(\omega; t) &= M_b(\omega) \beta_b(\omega) \gamma(\omega; t). \end{aligned}$$

Using the notations above, model (2.A.5) is the same as model (2.A.1) with latency $t_{0,k} = 0$, and there is no interaction term $P_k(\omega; t)Q_k(\omega; t)$.

2.B Lemmas, Theorems, and Their Proofs

In this section, we provide proofs of theorems and lemmas presented in our paper. Throughout this section, the time index collection $\mathcal{T}$ denotes $\{\tau \Delta\}_{\tau=0}^T$, and Ω is a discrete and finite set.

Lemma 2.B.1. *If $f : \mathcal{T} \rightarrow \mathbb{R}$ is not constant zero, its Fourier transform $\hat{f}(\xi)$ has at most finitely many zero points - $\xi \in \mathbb{R}$ such that $\hat{f}(\xi) = 0$ - in any compact subset of $\mathbb{R}$.*

Proof. Define the complex function $\Phi_f(z) := \frac{1}{T+1} \sum_{\tau=0}^T f(\tau\Delta) e^{-2\pi z(\tau\Delta)}$, for all $z = \eta + i\xi \in \mathbb{C}$. Since it is straightforward that $\Phi_f(z)$ satisfies the *Cauchy-Riemann* equation $\frac{\partial}{\partial \bar{z}} \Phi_f(z) = 0$, where $\frac{\partial}{\partial \bar{z}} = \frac{1}{2}(\frac{\partial}{\partial \eta} + i\frac{\partial}{\partial \xi})$ is a Wirtinger derivative, the Looman–Menchoff theorem implies that $\Phi_f(z)$ is a holomorphic function. Then the zero points of $\Phi_f(z)$ are isolated, i.e., every zero point has a neighbourhood that does not contain any other zero point (Theorem 3.7 of [Conway \(1973\)](#)). Therefore, $\hat{f}(\xi) = \Phi_f(i\xi)$ implies the desired result. $\square$

Theorem 2.B.1. *Suppose signals $\{Y_k(\omega; t) | t \in \mathcal{T}\}_{k=1}^K$, for $\omega \in \Omega$, are defined as in (2.A.1), $\mathbb{E}\beta_k = \mathbb{E}R_k(t) = 0$, and $t_{0,k} = \tau_{0,k}\Delta$ with some $\tau_{0,k} \in \mathbb{Z}$ for all $k = 1, 2, \dots, K$ and $t \in \mathcal{T}$. Let the random variable $U : \Omega \rightarrow \mathcal{T}$ be uniformly distributed on $\mathcal{T}$. Furthermore, we assume that $U(\omega)$, $\{\beta_k(\omega)\}_{k=1}^K$, and $\{R_k(\omega; t) | t \in \mathcal{T}\}_{k=1}^K$ are independent. Then, the autocovariance differences $\mathcal{A}_{kl}(s) := \mathbb{E}\{Y_k(t-U)Y_l(t+s-U)\} - \mathbb{E}\{R_k(t-U)R_l(t+s-U)\}$ depend only on s , the Fourier transforms of $\mathcal{A}(s)$ are*

$$\hat{\mathcal{A}}_{kl}(\xi) = \mathbb{E}(\beta_k \beta_l) \times \left| \hat{N}(\xi) \right|^2 \overline{\hat{h}_k(\xi)} \hat{h}_l(\xi) e^{2\pi i(t_{0,k} - t_{0,l})\xi}.$$

Furthermore, we have the following representation.

$$\mathcal{C}_{kl}(\xi) := |\hat{\mathcal{A}}_{kl}(\xi)| / \sqrt{|\hat{\mathcal{A}}_{kk}(\xi) \hat{\mathcal{A}}_{ll}(\xi)|} = |\text{corr}(\beta_k, \beta_l)|, \text{ for all } \xi \in \mathbb{R}.$$

Proof. The independence between $U(\omega)$, $\{\beta_k(\omega)\}_{k=1}^K$, and $\{R_k(\omega; t) | t \in \mathcal{T}\}_{k=1}^K$ implies

$$\begin{aligned} & \mathbb{E}[Y_k(t-U)Y_l(t+s-U)] - \mathbb{E}[R_k(t-U)R_l(t+s-U)] \\ &= \mathbb{E}(\beta_k \beta_l) \times \mathbb{E}[(N * h_k)(t - t_{0,k} - U) \times (N * h_l)(t + s - t_{0,l} - U)] \\ &= \mathbb{E}(\beta_k \beta_l) \times \frac{1}{T+1} \sum_{u=0}^T \left\{ (N * h_k)((\tau - \tau_{0,k} - u)\Delta) \times (N * h_l)((\underline{s} + \tau_{0,k} - \tau_{0,l})\Delta + (\tau - \tau_{0,k} - u)\Delta) \right\} \\ &= \mathbb{E}(\beta_k \beta_l) \times \frac{1}{T+1} \sum_{v=-(\tau - \tau_{0,k})}^{T - (\tau - \tau_{0,k})} \left\{ (N * h_k)(-v\Delta) \times (N * h_l)((\underline{s} + \tau_{0,k} - \tau_{0,l})\Delta - v\Delta) \right\} \\ &= \mathbb{E}(\beta_k \beta_l) \times [N * h_k(-\cdot)] * [N * h_l](s + (t_{0,k} - t_{0,l})), \end{aligned} \tag{2.B.1}$$

which only depends on s and does not depend on t . Here, the last equality follows from the periodic extension and the definition of convolution. Then, we have the following Fourier transform,

$$\begin{aligned} \hat{\mathcal{A}}_{kl}(\xi) &= \mathbb{E}(\beta_k \beta_l) \times \overline{(\widehat{N * h_k})(\xi)} (\widehat{N * h_l})(\xi) e^{2\pi i(t_{0,k} - t_{0,l})\xi} \\ &= \mathbb{E}(\beta_k \beta_l) \times \left| \hat{N}(\xi) \right|^2 \overline{\hat{h}_k(\xi)} \hat{h}_l(\xi) e^{2\pi i(t_{0,k} - t_{0,l})\xi}, \end{aligned}$$

for all $\xi \in \mathbb{R}$, which implies $\mathcal{C}_{kl}(\xi) = |\text{corr}(\beta_k, \beta_l)|$ for all $\xi \in \mathbb{R}$. $\square$

Lemma 2.B.2. *The vector-valued stochastic process $\mathbf{Z}(\omega; t) = \{(Z_1(\omega; t), Z_2(\omega; t), \dots, Z_K(\omega; t))^T\}_{t \in \mathcal{T}}$ is defined on probability space $(\Omega, 2^\Omega, \mathbb{P})$ and is weakly stationary with mean zero, $U : \Omega \rightarrow \mathcal{T}$ is uniformly distributed, and $U(\omega)$ and $\mathbf{Z}(\omega; t)$ are independent. Then the vector-valued stochastic process $\mathbf{Z}(\omega; t - U(\omega)) = \{(Z_1(\omega; t - U(\omega)), Z_2(\omega; t - U(\omega)), \dots, Z_K(\omega; t - U(\omega)))^T\}_{t \in \mathcal{T}}$ is weakly stationary with mean zero as well. Additionally, $\mathbb{E}[Z_k(t - U)Z_l(t + s - U)] = \mathbb{E}[Z_k(0)Z_l(s)]$, for all $k, l = 1, 2, \dots, K$.*

Proof. Let μ be the probability measure on $\mathcal{T}$ associated with the uniform random variable $U : \Omega \rightarrow \mathcal{T}$, i.e., $\mu = \mathbb{P} \circ U^{-1}$, then we have $\mathbb{E}\{Z_k(t - U)\} = \int_{\mathcal{T}} \mathbb{E}\{Z_k(t - U)|U = u\} \mu(du)$. The independence between U and $\mathbf{Z}$ implies $\mathbb{E}\{Z_k(t - U)|U = u\} = \mathbb{E}\{Z_k(t - u)\} = 0$ for all t , which results in $\mathbb{E}\{Z_k(t - U)\} = 0$ for all $t \in \mathcal{T}$. Similarly, we have

$$\begin{aligned} \mathbb{E}\{Z_k(t - U)Z_l(t + s - U)\} &= \int_{\mathcal{T}} \mathbb{E}\{Z_k(t - U)Z_l(t + s - U)|U = u\} \mu(du) \\ &= \int_{\mathcal{T}} \mathbb{E}\{Z_k(t - u)Z_l(t + s - u)\} \mu(du) \\ &= \int_{\mathcal{T}} \mathbb{E}\{Z_k(0)Z_l(s)\} \mu(du) \\ &= \mathbb{E}\{Z_k(0)Z_l(s)\}, \end{aligned}$$

This completes the proof of the desired result. $\square$

Theorem 2.B.2. *Suppose we have the following random vectors and stochastic processes.*

- $\mathfrak{B} := \{(\beta_k(\omega), \beta_l(\omega))\}_{\omega}^n$ are n i.i.d. random vectors.
- $\mathfrak{R} := \left\{ \left(\tilde{R}_k(\omega; t), \tilde{R}_l(\omega; t) \right) \mid t \in \mathcal{T} \right\}_{\omega=1}^n$ are n i.i.d. stochastic processes.
- $\mathfrak{W} := \{(W_k(\omega; t), W_l(\omega; t)) \mid t \in \mathcal{T}\}_{\omega=1}^n$ are n i.i.d. stochastic processes satisfying (i) random vectors $(W_k(\omega; t_1), W_l(\omega; t_1))$ and $(W_k(\omega; t_2), W_l(\omega; t_2))$ are independent whenever $t_1 \neq t_2$, (ii) the stochastic process $\{(W_k(\omega; t), W_l(\omega; t)) \mid t \in \mathcal{T}\}$ is weakly stationary with mean zero, and (iii) $\Sigma_{kl} = \mathbb{E}\{W_k(t)W_l(t)\}$ for all $t \in \mathcal{T}$.

The collections $\mathfrak{B}, \mathfrak{R}, \mathfrak{W}$ are independent. Furthermore, we have the following task-fMRI BOLD signals.[†]

$$Y_{k'}(\omega; t) = \beta_{k'}(\omega) \times N * h_{k'}(t - t_{0,k'}) + \tilde{R}_{k'}(\omega; t) + W_{k'}(\omega; t),$$

for $k' \in \{k, l\}$ and $\omega = 1, \dots, n$. We define the following estimator for $C_{kl}(\xi)$.

$$C_{kl}^{est,n}(\xi) := \left| \widehat{\mathcal{A}_{kl}^{est,n}}(\xi) \right| \bigg/ \sqrt{\left| \widehat{\mathcal{A}_{kk}^{est,n}}(\xi) \widehat{\mathcal{A}_{ll}^{est,n}}(\xi) \right|}, \quad \text{for all } \xi \in \mathbb{R},$$

[†]: Observed task-fMRI BOLD signals are $Y_k(\omega; t) = \beta_k(\omega) \times N * h_k(t - t_{0,k}) + R_k(\omega; t)$, where $R_k(\omega; t)$ are true reference signals. However, the true reference signals are not observable in applications and should be estimated by $\tilde{R}_k(\omega; t)$ using the AMUSE algorithm. The corresponding estimation bias $R_k(\omega; t) - \tilde{R}_k(\omega; t)$ is denoted by $W_k(\omega; t)$. Hence, we have $Y_k(\omega; t) = \beta_k(\omega) \times N * h_k(t - t_{0,k}) + \tilde{R}_k(\omega; t) + W_k(\omega; t)$, which is just a representation of $Y_k(\omega; t) = \beta_k(\omega) \times N * h_k(t - t_{0,k}) + R_k(\omega; t)$.

where $\widehat{(\cdot)}$ denotes the Fourier transform and

$$\begin{aligned} \mathcal{A}_{kl}^{est,n}(s) := & \frac{1}{T+1} \sum_{t \in \{\tau\Delta\}_{\tau=0}^T} \left[\frac{1}{n} \sum_{\omega=1}^n \{Y_k(\omega; t - U(\omega)) Y_l(\omega; t + s - U(\omega))\} \right] \\ & - \frac{1}{T+1} \sum_{t \in \{\tau\Delta\}_{\tau=0}^T} \left[\frac{1}{n} \sum_{\omega=1}^n \{\tilde{R}_k(\omega; t - U(\omega)) \tilde{R}_l(\omega; t + s - U(\omega))\} \right]. \end{aligned}$$

Then we have the following asymptotic behavior of $\mathcal{C}_{kl}^{est,n}(\xi)$ as $n \rightarrow \infty$.

$$\mathbb{P} \left\{ \lim_{n \rightarrow \infty} \mathcal{C}_{kl}^{est,n}(\xi) = \frac{\left| \mathbb{E}(\beta_k \beta_l) \frac{\widehat{\tilde{h}_k(\xi)} \widehat{\tilde{h}_l(\xi)}}{|\widehat{\tilde{h}_k(\xi)} \widehat{\tilde{h}_l(\xi)}|} e^{2\pi i \xi(t_{0,k} - t_{0,l})} + \frac{\Sigma_{kl}}{(T+1)|\widehat{\tilde{N}(\xi)}|^2 |\widehat{\tilde{h}_k(\xi)} \widehat{\tilde{h}_l(\xi)}|} \right|}{\sqrt{\left(\mathbb{E}(\beta_k^2) + \frac{\Sigma_{kk}}{(T+1)|\widehat{\tilde{N}(\xi)} \widehat{\tilde{h}_k(\xi)}|^2} \right) \left(\mathbb{E}(\beta_l^2) + \frac{\Sigma_{ll}}{(T+1)|\widehat{\tilde{N}(\xi)} \widehat{\tilde{h}_l(\xi)}|^2} \right)}}} \right\} = 1. \quad (2.B.2)$$

Proof. For each pair $(s, t) \in \mathcal{T}^2$, the strong law of large numbers implies that there exists $\mathcal{N}_{s,t} \in 2^\Omega$ depending on the pair (s, t) such that $\mathbb{P}(\mathcal{N}_{s,t}) = 0$ and

$$\begin{aligned} & \frac{1}{n} \sum_{\omega=1}^n \left[\{Y_k(\omega; t - U(\omega)) Y_l(\omega; t + s - U(\omega))\} - \{\tilde{R}_k(\omega; t - U(\omega)) \tilde{R}_l(\omega; t + s - U(\omega))\} \right] \\ & \rightarrow \mathbb{E} \{Y_k(t - U) Y_l(t + s - U)\} - \mathbb{E} \{\tilde{R}_k(t - U) \tilde{R}_l(t + s - U)\} =: \mathcal{D}_{kl}(t, s), \end{aligned}$$

in $\Omega - \mathcal{N}_{s,t}$ as $n \rightarrow \infty$. Then we have

$$\lim_{n \rightarrow \infty} \mathcal{A}_{kl}^{est,n}(s) = \frac{1}{T+1} \sum_{t \in \{\tau\Delta\}_{\tau=0}^T} \mathcal{D}_{kl}(t, s) \quad (2.B.3)$$

for all $s \in \mathcal{T}$ in $\Omega - \mathcal{N}$, where $\mathcal{N} := \bigcup_{s,t \in \mathcal{T}} \mathcal{N}_{s,t}$. The limit (2.B.3) implies the following in $\Omega - \mathcal{N}$.

$$\begin{aligned} \lim_{n \rightarrow \infty} \widehat{\mathcal{A}_{kl}^{est,n}}(\xi) &= \lim_{n \rightarrow \infty} \left\{ \frac{1}{T+1} \sum_{s \in \{\tau\Delta\}_{\tau=0}^T} \mathcal{A}_{kl}^{est,n}(s) e^{2\pi i \xi s} \right\} \\ &= \frac{1}{T+1} \sum_{s \in \{\tau\Delta\}_{\tau=0}^T} \lim_{n \rightarrow \infty} \{ \mathcal{A}_{kl}^{est,n}(s) \} e^{2\pi i \xi s} \\ &= \frac{1}{T+1} \sum_{t \in \{\tau\Delta\}_{\tau=0}^T} \left\{ \frac{1}{T+1} \sum_{s \in \{\tau\Delta\}_{\tau=0}^T} \mathcal{D}_{kl}(t, s) e^{2\pi i \xi s} \right\} \end{aligned}$$

for all $\xi \in \mathbb{R}$. In the derivation above, we change the order of summation and limit as the summation contains only finitely many terms. Since $\frac{1}{T+1} \sum_{s \in \{\tau\Delta\}_{\tau=0}^T} \mathcal{D}_{kl}(t, s) e^{2\pi i \xi s}$ is the Fourier transform of $\mathcal{D}_{kl}(t, s)$ with respect to s ,

repeating the calculation strategy in (2.B.1), we have

$$\frac{1}{T+1} \sum_{s \in \{\tau\Delta\}_{\tau=0}^T} \mathcal{D}_{kl}(t, s) e^{2\pi i \xi s} = \mathbb{E}(\beta_k \beta_l) \frac{\overline{\widehat{h}_k(\xi)} \widehat{h}_l(\xi)}{\left| \widehat{h}_k(\xi) \widehat{h}_l(\xi) \right|} e^{2\pi i \xi(t_{0,k} - t_{0,l})} + \frac{\Sigma_{kl}}{(T+1) \left| \widehat{N}(\xi) \right|^2 \left| \widehat{h}_k(\xi) \widehat{h}_l(\xi) \right|},$$

which does not depend on t . Then we have

$$\lim_{n \rightarrow \infty} \widehat{\mathcal{A}_{kl}^{est,n}}(\xi) = \mathbb{E}(\beta_k \beta_l) \frac{\overline{\widehat{h}_k(\xi)} \widehat{h}_l(\xi)}{\left| \widehat{h}_k(\xi) \widehat{h}_l(\xi) \right|} e^{2\pi i \xi(t_{0,k} - t_{0,l})} + \frac{\Sigma_{kl}}{(T+1) \left| \widehat{N}(\xi) \right|^2 \left| \widehat{h}_k(\xi) \widehat{h}_l(\xi) \right|} \quad (2.B.4)$$

for all $\xi \in \mathbb{R}$ in $\Omega - \mathcal{N}$. The limit (2.B.4) implies

$$\begin{aligned} \lim_{n \rightarrow \infty} \mathcal{C}_{kl}^{est,n}(\xi) &= \left| \lim_{n \rightarrow \infty} \widehat{\mathcal{A}_{kl}^{est,n}}(\xi) \right| / \sqrt{\left| \lim_{n \rightarrow \infty} \widehat{\mathcal{A}_{kk}^{est,n}}(\xi) \times \lim_{n \rightarrow \infty} \widehat{\mathcal{A}_{ll}^{est,n}}(\xi) \right|} \\ &= \frac{\left| \mathbb{E}(\beta_k \beta_l) \frac{\overline{\widehat{h}_k(\xi)} \widehat{h}_l(\xi)}{\left| \widehat{h}_k(\xi) \widehat{h}_l(\xi) \right|} e^{2\pi i \xi(t_{0,k} - t_{0,l})} + \frac{\Sigma_{kl}}{(T+1) \left| \widehat{N}(\xi) \right|^2 \left| \widehat{h}_k(\xi) \widehat{h}_l(\xi) \right|} \right|}{\sqrt{\left(\mathbb{E}(\beta_k^2) + \frac{\Sigma_{kk}}{(T+1) \left| \widehat{N}(\xi) \widehat{h}_k(\xi) \right|^2} \right) \left(\mathbb{E}(\beta_l^2) + \frac{\Sigma_{ll}}{(T+1) \left| \widehat{N}(\xi) \widehat{h}_l(\xi) \right|^2} \right)}} \end{aligned}$$

for all $\xi \in \mathbb{R}$ in $\Omega - \mathcal{N}$. Since $\mathbb{P}(\mathcal{N}) = \mathbb{P}\left(\bigcup_{s,t \in \mathcal{T}} \mathcal{N}_{s,t}\right) \leq \sum_{s,t \in \mathcal{T}} \mathbb{P}(\mathcal{N}_{s,t}) = 0$, the desired result (2.B.2) follows. $\square$

Proof of Theorem 2 in the manuscript: In the following equation provided in Theorem 1, $\mathbf{P}$ is a permutation matrix and $\mathbf{A}$ is a diagonal matrix.

$$\begin{pmatrix} 1 & 1 \\ \frac{1}{\beta^*} & 0 \end{pmatrix} = \mathbf{A} \mathbf{\Lambda}^{-1} \mathbf{P}^{-1}. \quad (2.B.5)$$

Then equation (2.B.5) implies $\beta^* = \beta_k(\omega) = a_{1i'}/a_{2i'}$ for some $i' \in \{1, 2\}$ and $a_{2j} = 0$ when $j \neq i'$. Furthermore, the pair $(\mathbf{A}, \mathbf{S}(\omega; t))$ also satisfies the following equation, which is provided in Theorem 1 as well.

$$\mathbf{A} \begin{pmatrix} s_1(\omega; t) \\ s_2(\omega; t) \end{pmatrix} = \begin{pmatrix} 1 & 1 \\ \frac{1}{\beta^*} & 0 \end{pmatrix} \begin{pmatrix} J_k(\omega; t) \\ R_k(\omega; t - U(\omega)) \end{pmatrix} = \begin{pmatrix} Y_k(\omega; t - U(\omega)) - \beta^* C_k \\ (N * h_k)(t - t_{0,k} - U(\omega)) - C_k \end{pmatrix}, \quad (2.B.6)$$

and equation (2.B.6) implies $a_{2i'} s_{i'}(\omega; t) = (N * h_k)(t - t_{0,k} - U(\omega)) - C_k$. Therefore, we have

$$a_{1i'} s_{i'}(\omega; t) = \frac{a_{1i'}}{a_{2i'}} \times a_{2i'} s_{i'}(\omega; t) = \beta_k(\omega) \times \{(N * h_k)(t - t_{0,k} - U(\omega)) - C_k\} = J_k(\omega; t).$$

2.C The Identifiability of Task-evoked Terms

Let $U : \Omega \rightarrow \mathcal{T}$ be a uniformly distributed random variable. Identifying the task-evoked terms $P_k(\omega; t)$ from $Y_k(\omega; t) = P_k(\omega; t) + R_k(\omega; t)$ is equivalent to identifying $P_k(\omega; t - U(\omega))$ from $Y_k(\omega; t - U(\omega))$ given the auxiliary random variable $U(\omega)$, which we can assumed to be known as it is artificially generated. Our proposed BOLD signal model $Y_k(\omega; t - U(\omega)) = P_k(\omega; t - U(\omega)) + R_k(\omega; t - U(\omega))$ can be represented in the following “mixing” form.

$$Y_k(\omega; t - U(\omega)) = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} \begin{pmatrix} P_k(\omega; t - U(\omega)) \\ R_k(\omega; t - U(\omega)) \end{pmatrix}, \quad (2.C.1)$$

where the identity matrix mixes the source signals $P_k(\omega; t - U(\omega))$ and $R_k(\omega; t - U(\omega))$.

In this section, we provide the identifiability of the task-evoked terms $P_k(\omega; t - U(\omega)) = \beta_k(\omega)N * h_k(t - t_{0,k} - U(\omega))$ as well as reference terms $R_k(\omega; t - U(\omega))$ in (2.C.1), up to deterministic coefficients, among the family of *mixing forms* defined as follows.

$$\mathcal{M} := \left\{ \begin{pmatrix} s_1(\omega; t) \\ s_2(\omega; t) \end{pmatrix} \middle| \exists \mathbf{A} \in \mathbb{R}^{2 \times 2} \text{ such that } Y_k(\omega; t - U(\omega)) = \mathbf{A} \begin{pmatrix} s_1(\omega; t) \\ s_2(\omega; t) \end{pmatrix} \right\}, \quad (2.C.2)$$

where all matrices $\mathbf{A}$ are deterministic 2-by-2 matrices mixing *source signals* $s_1(\omega; t)$ and $s_2(\omega; t)$. Specifically, we will show that, under some probabilistic conditions, all the stochastic signals $(s_1(\omega; t), s_2(\omega; t))$ in (2.C.2) are proportional to $(P_k(\omega; t - U(\omega)), R_k(\omega; t - U(\omega)))$ or $(R_k(\omega; t - U(\omega)), P_k(\omega; t - U(\omega)))$. Then the task-evoked terms $P_k(\omega; t - U(\omega))$ are identifiable in $\mathcal{M}$ in the following sense: there exist $i' \in \{1, 2\}$ and $\lambda \neq 0$ such that $s_{i'}(\omega; t) = \lambda P_k(\omega; t - U(\omega)) = \{\lambda \beta_k(\omega)\}N * h_k(t - t_{0,k} - U(\omega))$ for all $t \in \mathcal{T}$. Since the deterministic coefficient λ is canceled out in computing $\text{corr}(\beta_k, \beta_l)$, the identifiability of $P_k(\omega; t)$ up to a deterministic coefficient implies the exact identifiability of ptFC $|\text{corr}(\beta_k, \beta_l)|$.

The following theorem implies the identifiability of $P_k(\omega; t)$ and provides a formal foundation for the discussion above.

Theorem 2.C.1. *Let $U : \Omega \rightarrow \mathcal{T}$ be uniformly distributed. Suppose $U(\omega)$, $\beta_k(\omega)$, and $\{R_k(\omega; t) | t \in \mathcal{T}\}$ are independent, and $R_k(\omega; t)$ is WSMZ. Furthermore, we assume that there exists a $t^* \in \mathcal{T}$ such that*

$$\frac{\mathbb{E} \{P_k(t - U)P_k(t - t^* - U)\}}{\mathbb{E} \{P_k(t - U)\}^2} \neq \frac{\mathbb{E} \{R_k(t - U)R_k(t - t^* - U)\}}{\mathbb{E} \{R_k(t - U)\}^2}.$$

If there exists a WSMZ stochastic process $(s_1(\omega; t), s_2(\omega; t))$ such that $\{s_1(\omega; t)\}_{t \in \mathcal{T}}$ and $\{s_2(\omega; t)\}_{t \in \mathcal{T}}$ are uncorre-

lated,

$$\frac{\mathbb{E}\{s_1(t)s_1(t-t^*)\}}{\mathbb{E}\{s_1(t)\}^2} \neq \frac{\mathbb{E}\{s_2(t)s_2(t-t^*)\}}{\mathbb{E}\{s_2(t)\}^2}, \text{ and}$$

$$Y_k(\omega; t - U(\omega)) = P_k(\omega; t - U(\omega)) + R_k(\omega; t - U(\omega)) = \mathbf{A} \begin{pmatrix} s_1(\omega; t) \\ s_2(\omega; t) \end{pmatrix}$$

for some deterministic 2-by-2 matrix $\mathbf{A}$, then there exists a permutation matrix $\mathbf{P}$ and a non-singular diagonal matrix $\mathbf{\Lambda}$ such that

$$\begin{pmatrix} P_k(\omega; t - U(\omega)) \\ R_k(\omega; t - U(\omega)) \end{pmatrix} = \mathbf{P}\mathbf{\Lambda} \begin{pmatrix} s_1(\omega; t) \\ s_2(\omega; t) \end{pmatrix}. \quad (2.C.3)$$

Remark: (i) Our discussion of identifiability presented at the beginning of this section is based on (2.C.3). (ii) Since $\mathbb{E}\beta_k = \mathbb{E}R_k(t) = 0$ for all t , we have $\mathbb{E}\{P_k(t_1)R_k(t_2)\} = 0$ for all $t_1, t_2 \in \mathcal{T}$, i.e., $P_k(\omega; t_1)$ and $R_k(\omega; t_2)$ are uncorrelated. (iii) One can verify that $P_k(\omega; t - U(\omega)) = \beta_k(\omega)N * h_k(t - t_{0,k} - U(\omega))$ is WSMZ. (iv) Since $R_k(\omega; t)$ is WSMZ, Lemma 2.B.2 in the preceding section shows that $R_k(\omega; t - U(\omega))$ is WSMZ as well.

Based on these remarks, Theorem 2.C.1 is a straightforward result following from the Theorem 2 of Tong et al. (1991).

2.D The Performance of ptFCE in Estimating ptFCs

In this section, we analyze the performance of the ptFCE algorithm in terms of estimation bias and variance from both theoretical and simulational perspectives.

We first provide the estimation bias mechanism of the ptFCE algorithm motivated by

$$\mathbb{P} \left\{ \lim_{n \rightarrow \infty} C_{kl}^{est,n}(\xi) = \frac{\left| \mathbb{E}(\beta_k \beta_l) \frac{\widehat{h}_k(\xi) \widehat{h}_l(\xi)}{|\widehat{h}_k(\xi) \widehat{h}_l(\xi)|} e^{2\pi i \xi(t_{0,k} - t_{0,l})} + \frac{\Sigma_{kl}}{(T+1)|\widehat{N}(\xi)|^2 |\widehat{h}_k(\xi) \widehat{h}_l(\xi)|} \right|}{\sqrt{\left(\mathbb{E}(\beta_k^2) + \frac{\Sigma_{kk}}{(T+1)|\widehat{N}(\xi) \widehat{h}_k(\xi)|^2} \right) \left(\mathbb{E}(\beta_l^2) + \frac{\Sigma_{ll}}{(T+1)|\widehat{N}(\xi) \widehat{h}_l(\xi)|^2} \right)}}} \right\} = 1, \quad (2.D.1)$$

which is from Web Theorem 2.B.2. Because the quantities in

$$\Sigma_{kk} / [(T+1)|\widehat{N}(\xi) \widehat{h}_k(\xi)|^2] \approx 0, \quad \Sigma_{ll} / [(T+1)|\widehat{N}(\xi) \widehat{h}_l(\xi)|^2] \approx 0.$$

are not exactly zeros, the asymptotic estimation bias

$$\left| \lim_{n \rightarrow \infty} C_{kl}^{est,n}(\xi) - C_{kl}(\xi) \right| = \left| \lim_{n \rightarrow \infty} C_{kl}^{est,n}(\xi) - \frac{|\mathbb{E}(\beta_k \beta_l)|}{\sqrt{\mathbb{E}(\beta_k^2) \mathbb{E}(\beta_l^2)}} \right|$$

always exists. We further model that approximation residual $W_{k'}(\omega; t) = R_{k'}(\omega; t) - \tilde{R}_{k'}(\omega; t)$ at different nodes are independent, implying $\Sigma_{kl} = 0$ in (2.D.1). Then the following inequality, as a result of (2.D.1), indicates that our ptFCE algorithm tends to underestimate ptFC.

$$\begin{aligned} \lim_{n \rightarrow \infty} \mathcal{C}_{kl}^{est, n}(\xi) &=_{a.s.} \frac{|\mathbb{E}(\beta_k \beta_l)|}{\sqrt{\left(\mathbb{E}(\beta_k^2) + \frac{\Sigma_{kk}}{(T+1)|\hat{N}(\xi)\hat{h}_k(\xi)|^2} \right) \left(\mathbb{E}(\beta_l^2) + \frac{\Sigma_{ll}}{(T+1)|\hat{N}(\xi)\hat{h}_l(\xi)|^2} \right)}} \\ &\leq \frac{|\mathbb{E}(\beta_k \beta_l)|}{\sqrt{\mathbb{E}(\beta_k^2)\mathbb{E}(\beta_l^2)}} = |corr(\beta_k, \beta_l)|. \end{aligned} \quad (2.D.2)$$

Furthermore, the larger the variance of the noise $\epsilon_k(\omega; t)$ in

$$Y_k(\omega; t) = \beta_k(\omega) \times (N * h_k)(t - t_{0,k}) + R_k(\omega; t), \quad t \in \mathcal{T}, \quad k = 1, \dots, K, \quad \omega \in \Omega,$$

the less accurate is the approximation $R_{k'}(\omega; t) \approx \tilde{R}_{k'}(\omega; t)$ and the larger Σ_{kk} and Σ_{ll} in the inequality (2.D.2). Hence, the larger the variance of noise $\epsilon_k(\omega; t)$, the larger the underestimation bias. The formula above also indicates that the larger the underlying ptFC, the larger the underestimation bias. Subsequent simulations confirm all the conclusions on bias presented herein. We also show that the underestimation bias is moderate in general in terms of estimating ptFCs.

We present the performance of ptFCE in estimating ptFC by investigating estimation bias and variance using simulated data. Specifically, suppose the true ptFC in our simulation mechanism is ρ , and the estimated ptFC from this algorithm is $\hat{\rho}$, then we investigate bias $\hat{\rho} - \rho$ and the variance of $\hat{\rho}$. We implement the following data-generating mechanism, compatible with the definition of ptFC. The main steps of this mechanism are the following, while the details are provided in Appendix E.

Mechanism 0: **Step 1**, we generate random coefficients $(\beta_1(\omega), \beta_2(\omega))^T \sim N_2(\mathbf{0}, (\sigma_{ij})_{1 \leq i, j \leq 2})$ for $\omega = 1, \dots, n$, and $\rho := |\sigma_{12}|/\sqrt{\sigma_{11}\sigma_{22}}$ is the true underlying ptFC between synthetic nodes 1 and 2. **Step 2**, we generate reference signals $\{R_{k'}(\omega; \tau\Delta) = Q_{k'}(\omega; \tau\Delta) + \epsilon_{k'}(\omega; \tau\Delta)\}_{\tau=0}^T$ for $\omega = 1, \dots, n$ and $k' \in \{1, 2\}$, where the $2nT$ noise values $\{\epsilon_{k'}(\omega; \tau\Delta) | k' = 1, 2; \omega = 1, \dots, n; \tau = 0, \dots, T\} \sim_{iid} N(0, V)$. **Step 3**, we compute the signals $\{Y_{k'}(\omega; \tau\Delta) | k' = 1, 2\}_{\tau=0}^T$, for $\omega = 1, \dots, n$, by $Y_{k'}(\omega; \tau\Delta) = 9000 + \beta_{k'}(\omega) \times (N * h_{k'})(\tau\Delta) + R_{k'}(\omega; \tau\Delta)$.

We investigate sample sizes $n \in \{50, 100, 308, 1000\}$, where 308 is the sample size of the HCP dataset used to obtain the results in our paper. We apply the ptFCE algorithm to estimate the underlying ρ from synthetic signals. The corresponding estimate is denoted by $\hat{\rho}$. For each $\rho \in \{0.25, 0.5, 0.75\}$, we repeat this procedure 500 times. The resulting estimates $\hat{\rho}$ are summarized in Web Table 2.3 and Web Figure 2.4.

Additionally, we investigate the influence of random noise $\{\epsilon_{k'}(\omega; \tau\Delta)\}_{k', \tau, \omega} \sim_{iid} N(0, V)$ on the ptFCE algorithm. Specifically, we fix sample size $n = 308$; for each $\lambda \in [0.5, 5]$, we conduct the 500-run simulation study described above except we change the random noise to $\{\epsilon_{k'}(\omega; \tau\Delta)\}_{k', \tau, \omega} \sim_{iid} N(0, \lambda V)$. For each λ and

Table 2.3: The summaries of the estimated $\hat{\rho}$ in different underlying ptFC ρ scenarios. The percentage in the parenthesis after each mean shows the corresponding relative bias $(\hat{\rho} - \rho)/\rho$, where the minus signs indicate underestimation.

	$\rho = 0.25$		$\rho = 0.5$		$\rho = 0.75$	
	mean	sd	mean	sd	mean	sd
$n = 50$	0.255 (2%)	0.120	0.456 (-8.8%)	0.114	0.670 (-10.7%)	0.081
$n = 100$	0.246 (-1.6%)	0.093	0.460 (-8%)	0.079	0.671 (-10.5%)	0.055
$n = 308$	0.248 (-0.8%)	0.055	0.457 (-8.6%)	0.047	0.670 (-10.7%)	0.030
$n = 1000$	0.246 (-1.6%)	0.028	0.462 (-7.6%)	0.026	0.675 (-10%)	0.018

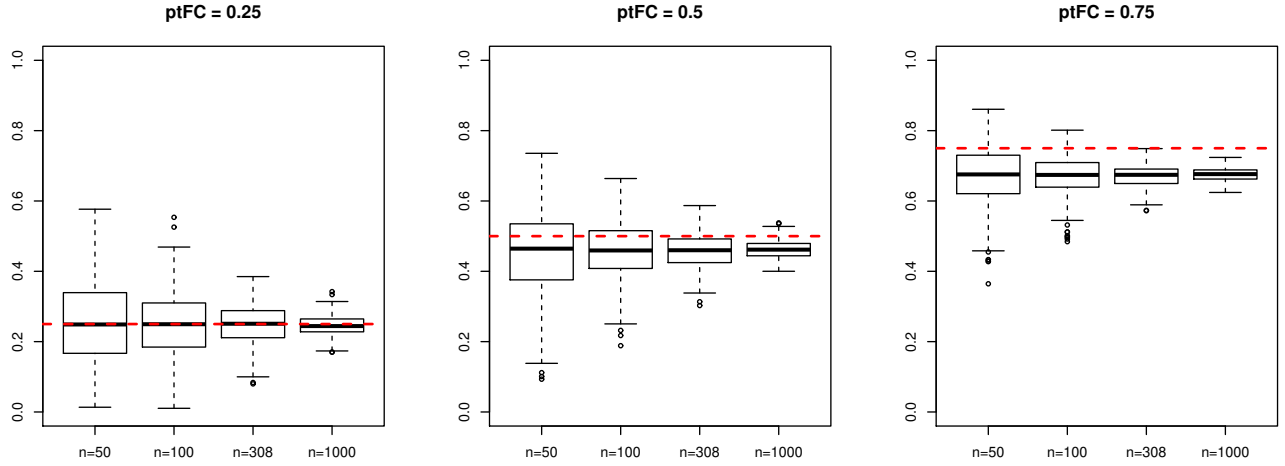


Figure 2.4: The boxplots summarizing the estimated $\hat{\rho}$ in different underlying ptFC ρ and different sample sizes n scenarios.

simulation run r , the estimated ptFC between synthetic nodes is denoted as $\hat{\rho}_{\lambda}^{(r)}$, then we obtain 500 curves $\{\hat{\rho}_{\lambda}^{(r)} | \lambda \in [0.5, 5]\}_{r=1}^{500}$ presented in Web Figure 2.5.

As expected from the theoretical analysis of the bias mechanism, our proposed ptFCE algorithm tends to underestimate true ptFCs, see Web Table 2.3 and Web Figures 2.4 and 2.5. But Web Table 2.3 shows that the bias is moderate. Simulations also confirm other conclusions on estimation bias presented in the theoretical analysis above. Additionally, Web Figure 2.3 shows that the magnitude of noise $\epsilon_k(\omega; t)$ does not influence the variance of the ptFCE estimates. The increase in sample size reduces estimation variance.

2.E Data-generating Mechanisms for Simulations

In this section, we provide the details of Mechanisms 0, 1, and 2 for generating synthetic signals in simulation studies in our paper. To mimic the HCP data of interest, we apply the following parameters in the two data-generating mechanisms.

- Repeation time $\Delta = 0.72$ (seconds)

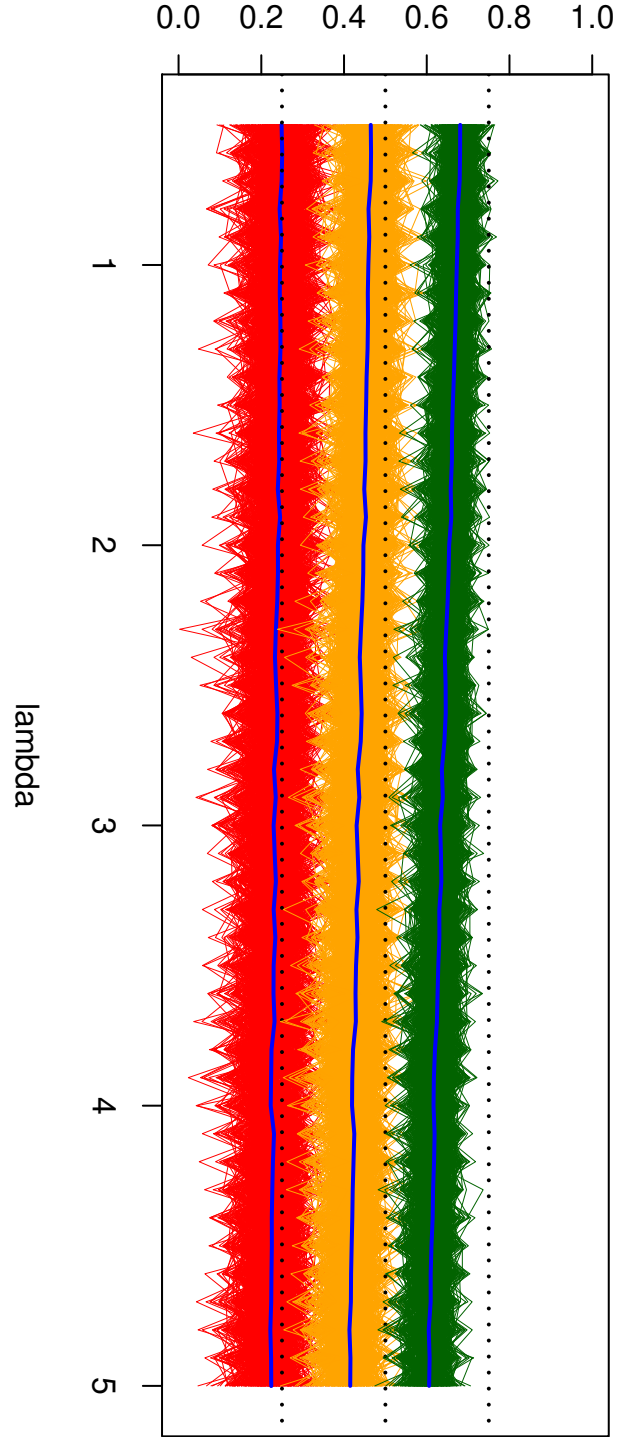


Figure 2.5: The red, orange, and green curves present the collections $\{\hat{\rho}_{\lambda}^{(r)} | \lambda \in [0.5, 5]\}_{r=1}^{500}$ corresponding to underlying ptFC $\rho = 0.25, 0.5, 0.75$. The three blue solid curves present the mean curves $\{\frac{1}{500} \sum_{r=1}^{500} \hat{\rho}_{\lambda}^{(r)} | \lambda \in [0.5, 5]\}$ in the three underlying ptFC scenarios, and the three black dotted lines present the three underlying ptFC.

- Number of observation time points $T = 283$.
- The task of interest (squeezing right toes) is presented by the stimulus signal $N(t) = \mathbf{1}_{[86.5, 98.5)}(t) + \mathbf{1}_{[162, 174)}(t)$.
- The tasks that are not of interest are presented by stimulus signals $\tilde{N}_1(t) = \mathbf{1}_{[71.35, 83.35)}(t) + \mathbf{1}_{[177.125, 189.125)}(t)$, $\tilde{N}_2(t) = \mathbf{1}_{[11, 23)}(t) + \mathbf{1}_{[116.63, 128.63)}(t)$, $\tilde{N}_3(t) = \mathbf{1}_{[26.13, 38.13)}(t) + \mathbf{1}_{[146.88, 158.88)}(t)$, and $\tilde{N}_4(t) = \mathbf{1}_{[56.26, 68.26)}(t) + \mathbf{1}_{[131.75, 143.75)}(t)$. They correspond to the tasks of squeezing left toes, squeezing left/right fingers, and moving tongue.
- HRF functions h_k and $\tilde{h}_{k,\gamma}$, for $k \in \{1, 3\}$ and $\gamma \in \{1, 2, 3, 4\}$, are the double-gamma variate functions implemented in the R function `canonicalHRF` with parameters `a1= 4`, `a2= 10`, `b1= 0.8`, `b2= 0.8`, `c= 0.4`; HRF functions h_2 and $\tilde{h}_{2,\gamma}$, for $\gamma \in \{1, 2, 3, 4\}$, are the same function with parameters `a1= 8`, `a2= 14`, `b1= 1`, `b2= 1`, `c= 0.3`. It is important to emphasize that none of the HRFs herein is canonical. The curves of all HRF functions used in our simulations studies are presented in Web Figure 2.6.

2.E.1 Mechanism 0

This mechanism is based on the task-fMRI BOLD signal model proposed in our paper. Specifically, we implement the following model to generate synthetic data.

$$Y_{k'}(\omega; t) = 9000 + \beta_{k'} \times N * h_{k'}(t) + \left\{ \sum_{\gamma=1}^4 \beta_{k',\gamma}(\omega) \times \tilde{N}_\gamma * \tilde{h}_{k',\gamma}(t) \right\} + \epsilon_{k'}(\omega; t), \quad (2.E.1)$$

where $t \in \mathcal{T} = \{\tau\Delta\}_{\tau=0}^T$ and $k' \in \{1, 2\}$.

We generate signals $\{(Y_1(\omega; \tau\Delta), Y_2(\omega; \tau\Delta))_{\tau=0}^T\}$, for $\omega = 1, \dots, n$, by the following steps.

- **Step 1:** Generate bivariate normal random vectors $(\beta_1(\omega), \beta_2(\omega))^T \sim_{i.i.d.} N_2((0, 0)^T, (\sigma_{ij})_{1 \leq i, j \leq 2})$ for $\omega = 1, \dots, n$, where $\rho := |\sigma_{12}|/\sqrt{\sigma_{11}\sigma_{22}}$ is the true underlying ptFC, and

$$\sigma_{11} = 2, \quad \sigma_{22} = 3, \quad \sigma_{12} = \rho \times \sqrt{\sigma_{11}\sigma_{22}}. \quad (2.E.2)$$

- **Step 2:** For each $\gamma \in \{1, 2, 3, 4\}$, generate bivariate normal random vector

$$\begin{pmatrix} \beta_{1,\gamma}(\omega) \\ \beta_{2,\gamma}(\omega) \end{pmatrix} \sim_{i.i.d.} N_2 \left(\begin{pmatrix} 0 \\ 0 \end{pmatrix}, \begin{pmatrix} 2 & 0.3 \times \sqrt{2 \times 3} \\ 0.3 \times \sqrt{2 \times 3} & 3 \end{pmatrix} \right),$$

for $\omega = 1, \dots, n$. The four collections $\{(\beta_{1,\gamma}(\omega), \beta_{2,\gamma}(\omega))^T\}_{\omega=1}^n$, for $\gamma \in \{1, 2, 3, 4\}$, are independently generated.

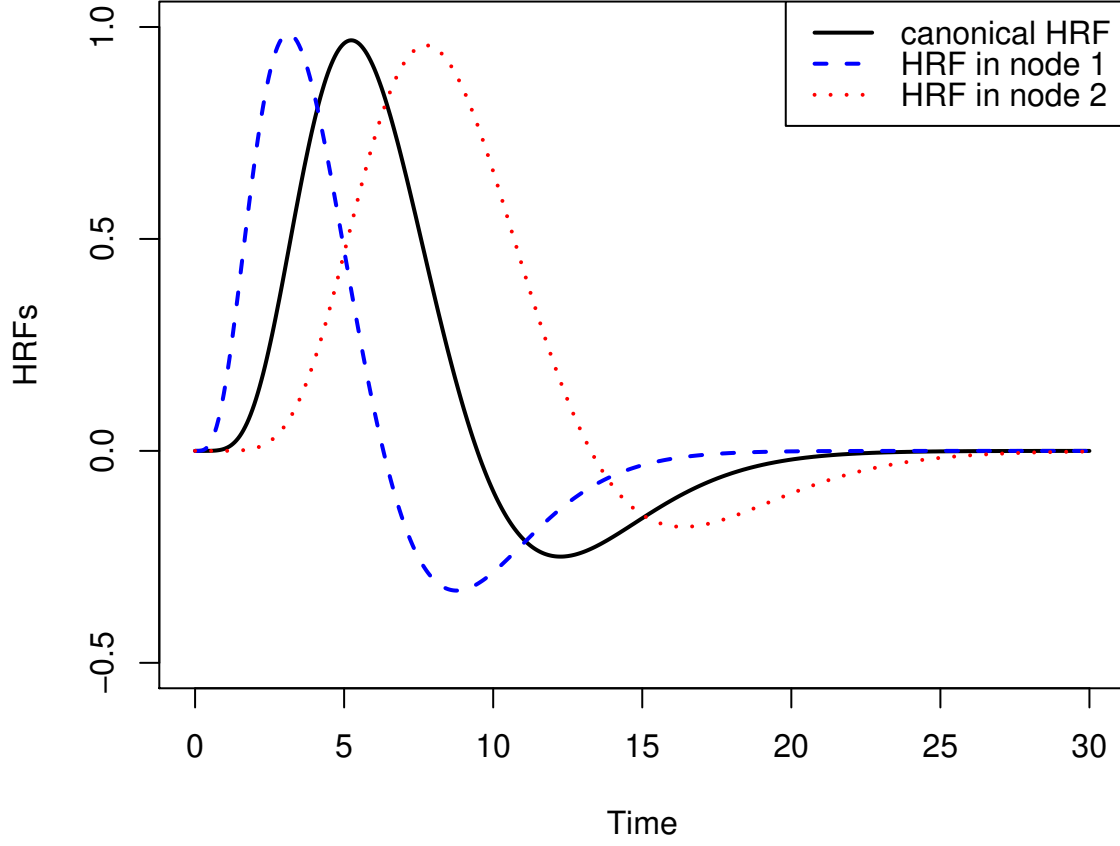


Figure 2.6: The black solid curve presents the canonical HRF (the R function `canonicalHRF` with default parameters). The blue dashed curve presents the R function `canonicalHRF` with parameters $a1=4$, $a2=10$, $b1=0.8$, $b2=0.8$, $c=0.4$. The red dotted curve presents the R function `canonicalHRF` with parameters $a1=8$, $a2=14$, $b1=1$, $b2=1$, $c=0.3$.

- **Step 3:** For each fixed $\omega \in \{1, \dots, n\}$, generate (white) noise as follows.

$$\begin{pmatrix} \epsilon_1(\omega; \tau\Delta) \\ \epsilon_2(\omega; \tau\Delta) \end{pmatrix} \sim_{i.i.d.} N_2 \left(\begin{pmatrix} 0 \\ 0 \end{pmatrix}, \begin{pmatrix} 30 & 0 \\ 0 & 30 \end{pmatrix} \right), \text{ for } \tau = 0, \dots, T.$$

The n collections $\{(\epsilon_1(\omega; \tau\Delta), \epsilon_2(\omega; \tau\Delta))\}_{\tau=0}^T$, for $\omega \in \{1, \dots, n\}$, are independently generated.

- **Step 4:** Compute the signals $\{(Y_1(\omega; \tau\Delta), Y_2(\omega; \tau\Delta))_{\tau=0}^T\}$, for $\omega = 1, \dots, n$, by model (2.E.1).

A pair of simulated $\{(Y_1(\omega; \tau\Delta), Y_2(\omega; \tau\Delta))_{\tau=0}^T\}$ using Mechanisms 0 is presented in Web Figure 2.7.

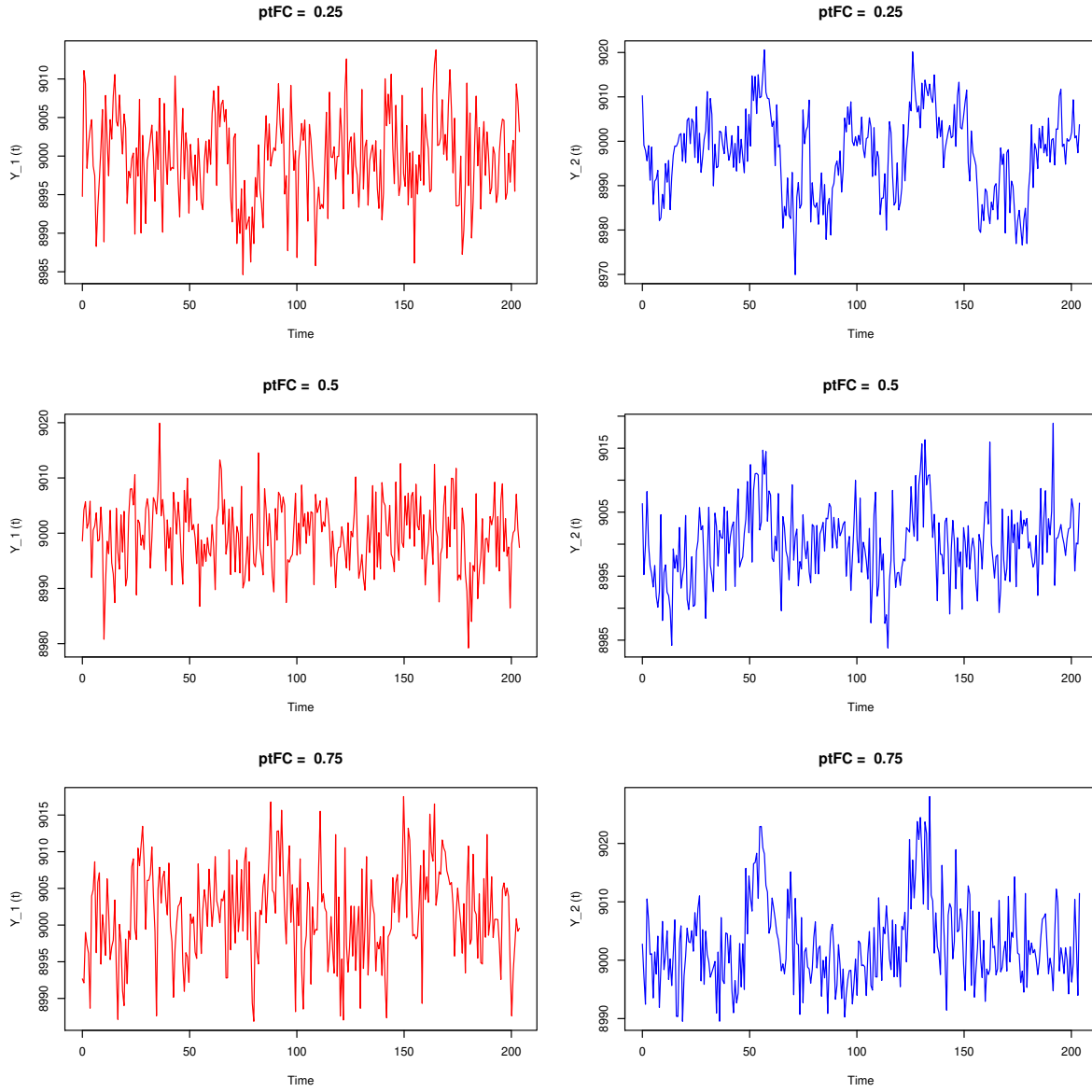


Figure 2.7: Left panels present signals generated for node 1, and right panels present signals generated for node 2. The underlying $ptFC$ s for pairs in the upper, middle, and lower rows are 0.25, 0.5, and 0.75, respectively.

2.E.2 Mechanism 1

Mechanism 1 is very similar to Mechanism 0 and is based on (2.E.1) as well, except for $k' \in \{1, 2, 3\}$. We generate signals $\{(Y_1(\omega; \tau\Delta), Y_2(\omega; \tau\Delta), Y_3(\omega; \tau\Delta))_{\tau=0}^T\}$, for $\omega = 1, \dots, n$, using the following steps.

- **Step 1:** Generate normal random vectors $(\beta_1(\omega), \beta_2(\omega), \beta_3(\omega))^T \sim_{i.i.d.} N_3(\mathbf{0}, (\sigma_{ij})_{1 \leq i, j \leq 3})$ for $\omega = 1, \dots, n$, where

$$(\sigma_{ij})_{1 \leq i, j \leq 3} = \begin{pmatrix} 2 & \rho_{12} \times \sqrt{2 \times 3} & 0 \\ \rho_{12} \times \sqrt{2 \times 3} & 3 & \rho_{23} \times \sqrt{2 \times 3} \\ 0 & \rho_{23} \times \sqrt{2 \times 3} & 2 \end{pmatrix}, \quad (2.E.3)$$

and $\rho_{ij} := |\sigma_{ij}| / \sqrt{\sigma_{ii}\sigma_{jj}}$ for $(i, j) \in \{(1, 2), (2, 3)\}$ are the true underlying ptFCs.

- **Step 2:** For each $\gamma \in \{1, 2, 3, 4\}$, generate bivariate normal random vector

$$\begin{pmatrix} \beta_{1,\gamma}(\omega) \\ \beta_{2,\gamma}(\omega) \\ \beta_{3,\gamma}(\omega) \end{pmatrix} \sim_{i.i.d.} N_3 \left(\begin{pmatrix} 0 \\ 0 \\ 0 \end{pmatrix}, \begin{pmatrix} 2 & 0.3 \times \sqrt{2 \times 3} & 0 \\ 0.3 \times \sqrt{2 \times 3} & 3 & 0.3 \times \sqrt{2 \times 3} \\ 0 & 0.3 \times \sqrt{2 \times 3} & 2 \end{pmatrix} \right),$$

for $\omega = 1, \dots, n$. The four collections $\{(\beta_{1,\gamma}(\omega), \beta_{2,\gamma}(\omega), \beta_{3,\gamma}(\omega))\}_{\omega=1}^n$, for $\gamma \in \{1, 2, 3, 4\}$, are independently generated.

- **Step 3:** For each fixed $\omega \in \{1, \dots, n\}$, generate (white) noise as follows.

$$\begin{pmatrix} \epsilon_1(\omega; \tau\Delta) \\ \epsilon_2(\omega; \tau\Delta) \\ \epsilon_3(\omega; \tau\Delta) \end{pmatrix} \sim_{i.i.d.} N_3 \left(\begin{pmatrix} 0 \\ 0 \\ 0 \end{pmatrix}, \begin{pmatrix} 30 & 0 & 0 \\ 0 & 30 & 0 \\ 0 & 0 & 30 \end{pmatrix} \right), \text{ for } \tau = 0, \dots, T.$$

The n collections $\{(\epsilon_1(\omega; \tau\Delta), \epsilon_2(\omega; \tau\Delta), \epsilon_3(\omega; \tau\Delta))_{\tau=0}^T\}$, for $\omega \in \{1, \dots, n\}$, are independently generated.

- **Step 4:** Compute the signals $\{(Y_1(\omega; \tau\Delta), Y_2(\omega; \tau\Delta), Y_3(\omega; \tau\Delta))_{\tau=0}^T\}$, for $\omega = 1, \dots, n$, by model (2.E.1).

2.E.3 Mechanism 2

This mechanism is motivated by the Pearson correlation approach and implemented by the following steps.

- **Step 1:** For each fixed $\omega \in \{1, \dots, n\}$, we independently generate $(\epsilon_1(\omega; \tau\Delta), \epsilon_2(\omega; \tau\Delta), \epsilon_3(\omega; \tau\Delta))^T$, for $\tau \in \{1, \dots, T\}$, using the following distributions:

(i) If $N(\tau\Delta) = 1$, we apply

$$\begin{pmatrix} \epsilon_1(\omega; \tau\Delta) \\ \epsilon_2(\omega; \tau\Delta) \\ \epsilon_3(\omega; \tau\Delta) \end{pmatrix} \sim N_3 \left(\begin{pmatrix} 0 \\ 0 \\ 0 \end{pmatrix}, \begin{pmatrix} 30 & \varrho \times 30 & 0 \\ \varrho \times 30 & 30 & \varrho \times 30 \\ 0 & \varrho \times 30 & 30 \end{pmatrix} \right); \quad (2.E.4)$$

if $N(\tau\Delta) = 0$, we implement

$$\begin{pmatrix} \epsilon_1(\omega; \tau\Delta) \\ \epsilon_2(\omega; \tau\Delta) \\ \epsilon_3(\omega; \tau\Delta) \end{pmatrix} \sim N_3 \left(\begin{pmatrix} 0 \\ 0 \\ 0 \end{pmatrix}, \begin{pmatrix} 30 & 0 & 0 \\ 0 & 30 & 0 \\ 0 & 0 & 30 \end{pmatrix} \right).$$

The n collections $\{(\epsilon_1(\omega; \tau\Delta), \epsilon_2(\omega; \tau\Delta), \epsilon_3(\omega; \tau\Delta))\}_{\tau=0}^T$, for $\omega \in \{1, \dots, n\}$, are independently generated.

- **Step 2:** Compute the signals $\{(Y_1(\omega; \tau\Delta), Y_2(\omega; \tau\Delta), Y_3(\omega; \tau\Delta))\}_{\tau=0}^T$, for $\omega = 1, \dots, n$, by the following.

$$Y_{k'}(\omega; t) = 9000 + N * h_{k'}(t) + \left\{ \sum_{\gamma=1}^4 \tilde{N}_\gamma * \tilde{h}_{k',\gamma}(t) \right\} + \epsilon_{k'}(\omega; t), \quad (2.E.5)$$

where $t \in \mathcal{T} = \{\tau\Delta\}_{\tau=0}^T$ and $k' \in \{1, 2, 3\}$.

Simulated $\{(Y_1(\omega; \tau\Delta), Y_2(\omega; \tau\Delta), Y_3(\omega; \tau\Delta))\}_{\tau=0}^T$ using Mechanisms 1 and 2, are presented in Web Figure 2.8.

2.F Limitations of the Pearson correlation approach

Web Figure 2.9 illustrates the limitations of the Pearson correlation approach defined as follows in terms of the influence of latency/HRF variance and noise.

$$|corr(P_k, P_l)| = \left| \int_{\mathcal{T}} \phi_k(t) \times \phi_l(t) \mu(dt) \right| / \sqrt{\int_{\mathcal{T}} |\phi_k(t)|^2 \mu(dt) \times \int_{\mathcal{T}} |\phi_l(t)|^2 \mu(dt)}. \quad (2.F.1)$$

2.G Details of Beta-series Regression and Coherence Analysis

Beta-series regression: We apply the procedure described in Chapter 9.2 of [Ashby \(2019\)](#) to estimate task-evoked FC using beta-series regression, except we ignore the “nuisance term” therein. For each subject ω and underlying ρ_{ij} or ϱ_{ij} with $i < j$, we apply beta-series regression to signals $\{Y_i(\omega; \tau\Delta), Y_j(\omega; \tau\Delta)\}_{\tau=0}^T$, estimate the FC evoked by the task of interest $N(t)$, and denote the estimated quantity by $\hat{\rho}_{ij,\omega}^{betaS}$. Then we compute the mean and median of $\{\hat{\rho}_{ij,\omega}^{betaS}\}_{\omega=1}^n$ across all $\omega = 1, \dots, n$ and denote them by $\hat{\rho}_{ij,mean}^{betaS}$ and $\hat{\rho}_{ij,median}^{betaS}$.

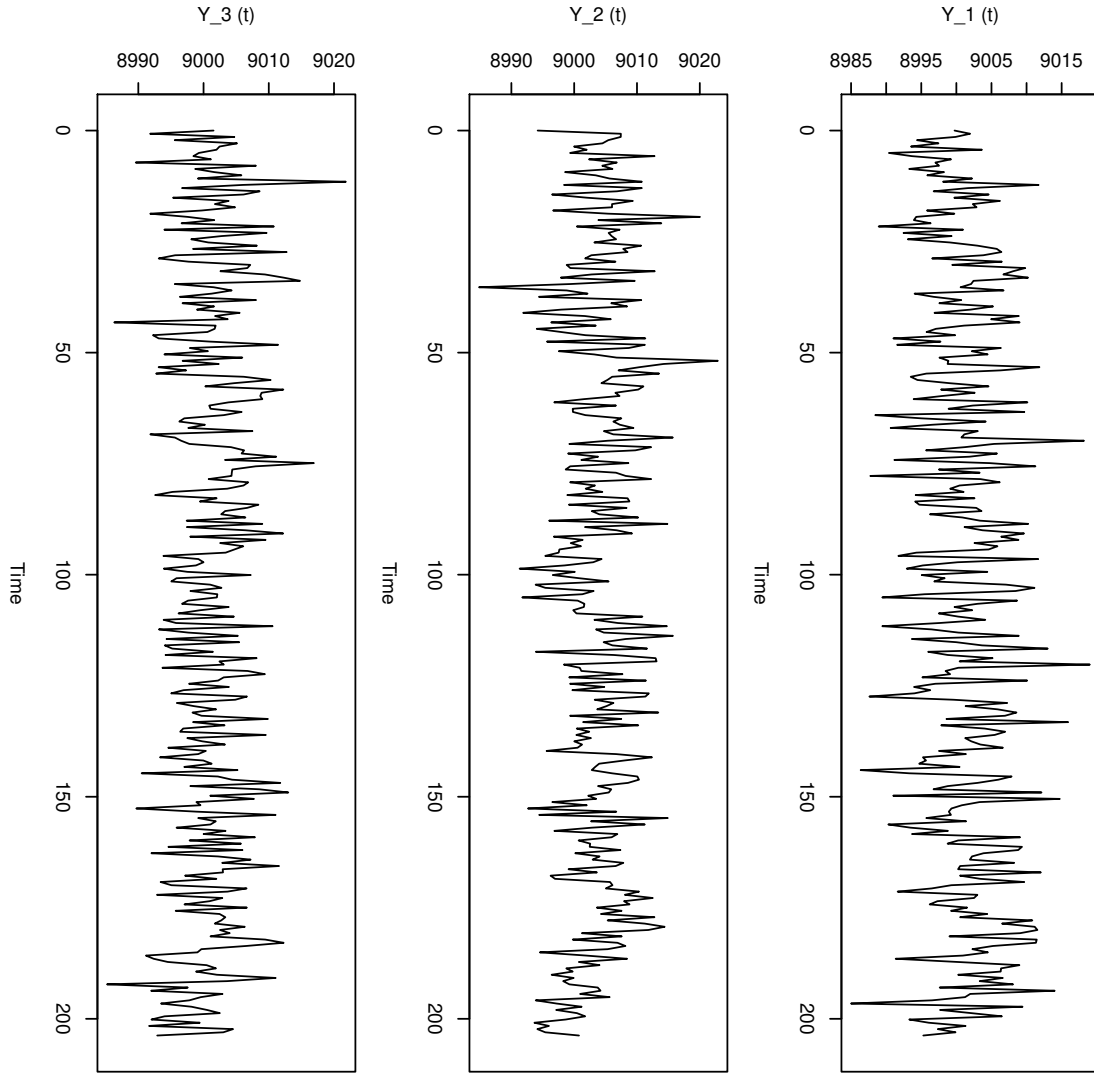


Figure 2.8: Signals generated for nodes 1, 2, and 3 using Mechanism 2.

Coherence analysis: For each ω and ρ_{ij} or ϱ_{ij} , compute the coherence between signals $\{Y_i(\omega; \tau\Delta)\}_\tau$ and $\{Y_j(\omega; \tau\Delta)\}_\tau$ by the R function `coh` in package `seewave`. The coherence is a function $coh_{ij,\omega}(\xi)$ of ξ . Since HRFs act as band-pass filters (0 – 0.15 Hz, [Aguirre et al. \(1997\)](#)), compute the median of $coh_{ij,\omega}(\xi)$ across all $\xi \in (0, 0.15)$ and denote it by $\hat{\rho}_{ij,\omega}^{Coh}$. The mean and median of $\{\hat{\rho}_{ij,\omega}^{Coh}\}_{\omega=1}^n$ across all ω are denoted by $\hat{\rho}_{ij,mean}^{Coh}$ and $\hat{\rho}_{ij,median}^{Coh}$, respectively.

2.H Future Research

An interesting extension of our proposed model is the inclusion of an interaction term between the $P_k(\omega; t)$ and $Q_k(\omega; t)$ in model $Y_k(\omega; t) = P_k(\omega; t) + Q_k(\omega; t) + \epsilon_k(\omega; t)$. In order to experimentally validate whether the inclusion of such an interaction term would be biologically relevant, novel experimental designs would be needed. Particularly, we may consider an experiment where the subjects are at rest for a certain period during the scanning session

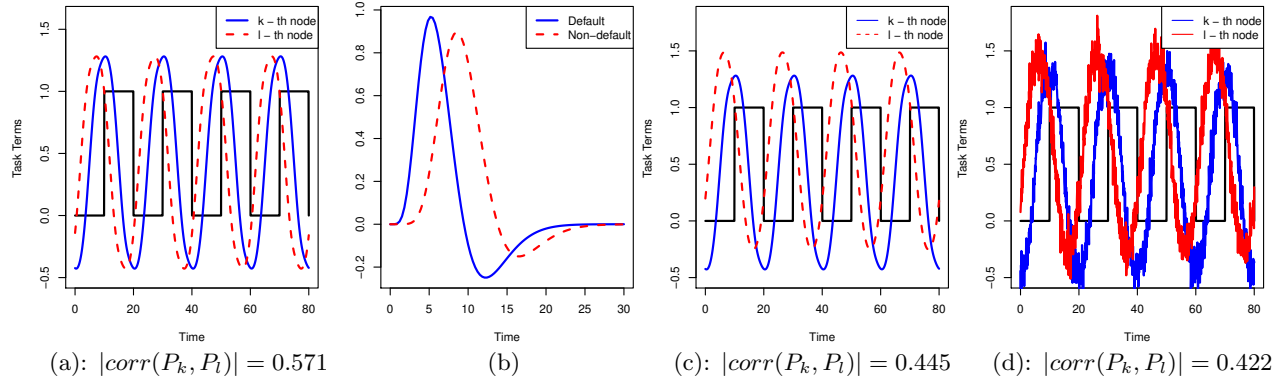


Figure 2.9: An example illustrating the limitations of the Pearson correlation approach in (2.F.1). Let $N(t) = \sum_{m=1}^4 \mathbf{1}_{(20m-10, 20m)}(t)$ be the stimulus signal of a block design, $\mathcal{T} = [0, 80]$, and h_k be the `canonicalHRF` function with default parameters in the R package `neuRosim`. h_k is illustrated by the solid blue curve in (b). Panel (a) shows the influence of the variation in latency on (2.F.1), where $h_l(t) = h_k(t + 3)$, i.e., $t_{0,k} = -3$; the task-evoked terms at the k^{th} and l^{th} nodes are presented by blue and red curves, respectively. In (b), h_l is replaced by the `canonicalHRF` function with parameters $\mathbf{a1} = 10, \mathbf{a2} = 15, \mathbf{b1} = \mathbf{b2} = 0.9, \mathbf{c} = 0.35$, and presented by the dashed red curve. Panel (c) shows the influence of variation in HRF on (2.F.1), and the task-evoked terms at the k^{th} and l^{th} nodes are presented by blue and red curves, respectively, where h_l is defined in (b) as the red curve. Panel (d) is a noise-contaminated version of (c), i.e., curves of $P_k(t) + \epsilon_k(t)$ (blue) and $P_l(t) + \epsilon_l(t)$ (red) with $\epsilon_k(t), \epsilon_l(t) \sim_{iid} N(0, 0.1)$ for each t . Panel (d) shows that random noise can further influence (2.F.1). In panels (a,c,d), the presented $|corr(P_k, P_l)|$ is computed by (2.F.1).

followed by the performance of the task. Using this design, we may obtain estimates of P_k and Q_k considering the model with and without an interaction term and discuss the biological relevance of the results. Future work may investigate this issue and develop theoretical conditions for the identifiability of terms in a model with an interaction term and estimation algorithms.

In many applications, task-evoked FC at the subject level instead of the population level is of interest. In our subsequent research, we will propose a framework parallel to the ptFC one at the subject level. Additionally, implementing the *persistent homology* (PH) framework to FC estimates is an effective way of circumventing the choice of threshold for FC measurements, e.g., see Lee et al. (2001). Applying the PH approach to our proposed ptFCs is left for future research as well.

Chapter 3

Randomness and Statistical Inference of Shapes via the Smooth Euler Characteristic Transform

3.1 Introduction

The quantification of shapes has become an important research direction in both the applied and theoretical sciences. It has brought advances to many fields including network analysis ([Lee et al., 2011](#)), geometric morphometrics ([Boyer et al., 2011](#); [Gao et al., 2019](#)), biophysics and structural biology ([Wang et al., 2021](#)), and radiomics (i.e., the field that aims to correlate clinical imaging features and genomic assays) ([Crawford et al., 2020](#)). If shapes are considered random, then their corresponding quantitative summaries are also random — implying such summaries of random shapes are statistics. The statistical inference of shapes based on these quantitative summaries has been of particular interest (e.g., the derivation of confidence sets by [Fasy et al. \(2014\)](#) and a framework for hypothesis testing by [Robinson and Turner \(2017\)](#)).

-
- A version of this chapter appeared in [Meng et al. \(2022a\)](#).
 - **Author contributions:** Kun Meng developed the theoretical framework of this chapter, proved the theorems, and conducted simulation studies. Ani Eloyan and Lorin Crawford supervised this chapter. Kun Meng, Lorin Crawford, and Ani Eloyan wrote version [Meng et al. \(2022a\)](#). We — Kun Meng, Lorin Crawford, and Ani Eloyan — want to thank Dr. Matthew T. Harrison for the suggestion to implement the permutation test. This suggestion greatly improved the quality of this chapter.
 - **Abbreviations:** ECT, Euler characteristic transform; ECC, Euler characteristic curve; GBM, glioblastoma multiforme; HCP, homotopy critical point; i.i.d., independent and identically distributed; LECT, lifted Euler characteristic transform; PD, persistent diagram; PHT, persistent homology transform; PECT, primitive Euler characteristic transform; RKHS, reproducing kernel Hilbert space; RCLL, right continuous with left limits; SECT, smooth Euler characteristic transform; TDA, Topological data analysis.

3.1.1 Overview of Topological Data Analysis

Topological data analysis (TDA) presents a collection of statistical methods that quantitatively summarize the shapes represented in data using computational topology (Edelsbrunner and Harer, 2010). One common statistical invariant in TDA is the persistent diagram (PD) (Edelsbrunner et al., 2000). When equipped with the p -th Wasserstein distance for $1 \leq p < \infty$, the collection of PDs denoted as $\mathcal{D}$ is a Polish space (Mileyko et al., 2011); hence, probability measures can be applied, and the randomness of shapes can be potentially represented using the probability measures on $\mathcal{D}$. However, a single PD does not preserve all relevant information of a shape (Crawford et al., 2020). Using the Euler calculus, Ghrist et al. (2018) (Corollary 6 therein) showed that the persistent homology transform (PHT) (Turner et al., 2014b), motivated by integral geometry and differential topology, concisely summarizes information within shapes. The PHT takes values in $C(\mathbb{S}^{d-1}; \mathcal{D}^d) = \{\text{all continuous maps } F : \mathbb{S}^{d-1} \rightarrow \mathcal{D}^d\}$, where $\mathbb{S}^{d-1}$ denotes the sphere $\{x \in \mathbb{R}^d : \|x\| = 1\}$ and $\mathcal{D}^d$ is the d -fold Cartesian product of $\mathcal{D}$ (see Lemma 2.1 and Definition 2.1 of Turner et al. (2014b)). Since $\mathcal{D}$ is not a vector space and the distances on $\mathcal{D}$ (e.g., the p -th Wasserstein and bottleneck distances (Cohen-Steiner et al., 2007)) are very abstract, many fundamental concepts in classical statistics are not easy to implement with summaries resulting from the PHT. For example, the definition of moments corresponding to probability measures on $\mathcal{D}$ (e.g., means and variances) is highly nontrivial (see Mileyko et al. (2011) and Turner et al. (2014a)). The difficulty in defining these properties prevents the applications of PHT-based statistical inference methods in $C(\mathbb{S}^{d-1}; \mathcal{D}^d)$.

The smooth Euler characteristic transform (SECT) proposed by Crawford et al. (2020) provides an alternative summary statistic for shapes. The SECT not only preserves the information of shapes of interest (see Corollary 6 of Ghrist et al. (2018)), but it also represents shapes using univariate continuous functions instead of PDs. Specifically, values of the SECT are maps from the sphere $\mathbb{S}^{d-1}$ to a separable Banach space — the collection of real-valued continuous functions on a compact interval, say $C([0, T]) = \mathcal{B}$ for some $T > 0$ (values of T will be given in Eq. (3.2.1)). That is, for any shape K , its SECT, denoted as $SECT(K) = \{SECT(K)(\nu)\}_{\nu \in \mathbb{S}^{d-1}}$, is in $\mathcal{B}^{\mathbb{S}^{d-1}} = \{\text{all maps } f : \mathbb{S}^{d-1} \rightarrow \mathcal{B}\}$; specifically, $SECT(K)(\nu) \in \mathcal{B}$ for each $\nu \in \mathbb{S}^{d-1}$. Therefore, the randomness of shapes K is represented via the SECT by a collection of $\mathcal{B}$ -valued random variables. The probability theory on separable Banach spaces has been better developed than on $\mathcal{D}$. Specifically, a $\mathcal{B}$ -valued random variable is a stochastic process with its paths in $\mathcal{B}$ (we will further show in Section 3.2.1 that $\mathcal{B}$ herein can be replaced with a reproducing kernel Hilbert space (RKHS)), and the theory of stochastic processes has been developed for nearly a century. Many mathematical tools are available for building the foundations of the randomness and statistical inference of shapes.

The work from Crawford et al. (2020) applied the SECT to magnetic resonance images taken from tumors of a cohort of glioblastoma multiforme (GBM) patients. Using summary statistics derived from the SECT as predictors of the Gaussian process regression, the authors showed that the SECT has the power to predict clinical outcomes better than existing tumor shape quantification approaches and common molecular assays. The relative

performance of the SECT in the GBM study represents a promising future for the utility of SECT in medical imaging and more general statistical applications investigating shapes. Similarly, Wang et al. (2021) applied derivatives of the Euler characteristic transform as predictors of statistical models for subimage analysis which is related to the task of variable selection and seeks to identify physical features that are most important for differentiating between two classes of shapes. In both Crawford et al. (2020) and Wang et al. (2021), the shapes were implemented as the predictors of regressions, and the randomness of these predictors was ignored.

3.1.2 Overview of Contributions

In this paper, we model the distributions of shapes via the SECT using RKHS-valued random fields and provide the corresponding foundations using tools in algebraic and computational topology, Sobolev spaces, and functional analysis. In contrast to work like Crawford et al. (2020) and Wang et al. (2021), we model realizations from the SECT as the responses of an RKHS-valued random field rather than as input variables or predictors. Modeling the distributions of shapes helps answer the following statistical inference question: *Is the difference between two groups of shapes significant or just random?* Based on the foundations of the randomness of shapes, we propose an approach for testing hypotheses on random shapes and answer this statistical inference question.

Using the homotopy theory and PDs, we first propose a general collection of shapes on which the SECT is well-defined. Then, we show that, for each shape in this collection, its SECT is in $C(\mathbb{S}^{d-1}; \mathcal{H}) = \{\text{all continuous maps } F : \mathbb{S}^{d-1} \rightarrow \mathcal{H}\}$, where $\mathcal{H} = H_0^1([0, T])$ is not only a Sobolev space (see Chapter 8.3 of Brezis (2011)) but also an RKHS[†]. We further construct a probability space to model the distributions of shapes. This probability space makes SECT an $\mathcal{H}$ -valued random field indexed by $\mathbb{S}^{d-1}$ with continuous paths, which is equivalent to that SECT being a $C(\mathbb{S}^{d-1}; \mathcal{H})$ -valued random variable. Additionally, the restriction of the SECT to any point on $\mathbb{S}^{d-1}$ is a real-valued stochastic process with paths in $\mathcal{H}$. Based on the proposed probability space, we provide an assumption on the moments of the SECT and define the form of its mean and covariance. Using a Sobolev embedding result and functional analysis arguments, we show some properties of the mean and covariance for the SECT. These properties allow us to develop the Karhunen–Loève expansion of the SECT, which is the key tool for our proposed hypothesis testing framework.

Traditionally, the statistical inference of shapes in TDA is conducted in persistent diagram space $\mathcal{D}$, which is not suitable for exponential family-based distributions and requires any corresponding statistical inference to be highly nonparametric. For example, Fasy et al. (2014) used subsampling, and Robinson and Turner (2017) applied permutation tests. The PHT-based statistical inference in $C(\mathbb{S}^{d-1}; \mathcal{D}^d)$ is even more difficult. With the

†: Strictly speaking, the functions in Sobolev space $\mathcal{H}$ are defined on the open interval $(0, T)$ instead of the closed interval $[0, T]$ (see Chapter 8.2 of Brezis (2011)). Hence, the rigorous notation of $\mathcal{H}$ should be $H_0^1((0, T))$. However, Theorem 8.8 of Brezis (2011) indicates that each function in $H_0^1((0, T))$ can be uniquely represented by a continuous function defined on $[0, T]$, which implies that functions in $H_0^1((0, T))$ can be viewed as defined on the closed interval $[0, T]$. Therefore, to implement the boundary values on $\partial(0, T) = \{0, T\}$, we use the notation $H_0^1([0, T])$ throughout this paper to indicate that all functions in $\mathcal{H}$ are viewed as defined on $[0, T]$. The same reasoning is applied for the space $W_0^{1,p}([0, T])$ implemented later in this paper (see Theorem 8.8 and the Remark 8 after Proposition 8.3 in Brezis (2011)).

Karhunen–Loève expansion of the SECT and the central limit theorem, some hypotheses on the distributions of shapes can be tested using the normal distribution-based methods (e.g., the adaptive Neyman test (Fan, 1996) and the χ^2 -test). We will provide the details of these methods in Section 3.4.

One can use a morphology example from Turner et al. (2014b) to illustrate a potential application of SECT-based hypothesis testing. The study of variation of complex traits and phenotypes is fundamental to understanding hypotheses in evolutionary biology. Suppose we observe heel bones from two groups of primates (i.e., heel bones $\{K_i^{(j)}\}_{i=1}^n$ are from the j -th group for $j \in \{1, 2\}$). For each group j , suppose the underlying distribution of statistics $\{SECT(K_i^{(j)})\}_{i=1}^n$ is $\mathbf{P}^{(j)}$, and $\mathbf{P}^{(j)}$ is a probability measure on space $C(\mathbb{S}^{d-1}; \mathcal{H})$. Then, assessing the evolutionary biological question of “whether the two groups of primates are of different species” may be boiled down to testing the hypotheses H_0 : “ $\mathbf{P}^{(1)}$ and $\mathbf{P}^{(2)}$ have the same mean” vs. H_1 : “ $\mathbf{P}^{(1)}$ and $\mathbf{P}^{(2)}$ do not have the same mean” (the rigorous definition of “the mean of $\mathbf{P}^{(j)}$ ” and the rigorous version of the hypotheses will be provided in Section 3.3.2 and Eq. (3.4.1), respectively, of this paper).

3.1.3 Relevant Notation and Paper Organization

Throughout this paper, a “shape” refers to a triangulable subset of $\mathbb{R}^d$ defined as follows.

Definition 3.1.1. (i) Let K be a subset of $\mathbb{R}^d$. If there exists a finite simplicial complex S such that the corresponding polygon $|S| = \bigcup_{s \in S} s$ is homeomorphic to K , we say K is triangulable. (ii) Let $\mathcal{S}_d$ denote the collection of all triangulable subsets of $\mathbb{R}^d$.

The definitions of simplicial complexes and triangulable spaces can be found in Munkres (2018). The triangulability assumption is standard in the TDA literature (e.g., Cohen-Steiner et al. (2010) and Turner et al. (2014b)). Throughout this paper, we apply the following:

- i) All the linear spaces are defined with respect to field $\mathbb{R}$. For any normed space $\mathcal{V}$, let $\|\cdot\|_{\mathcal{V}}$ denote its norm. Let $\|x\|$ denote the Euclidean norm if x is a finite-dimensional vector.
- ii) Let X be a compact metric space equipped with metric d_X and let $\mathcal{V}$ denote a normed space. $C(X; \mathcal{V})$ is the collection of all continuous maps from X to $\mathcal{V}$. Furthermore, $C(X; \mathcal{V})$ is a normed space equipped with the norm

$$\|f\|_{C(X; \mathcal{V})} = \sup_{x \in X} \|f(x)\|_{\mathcal{V}}.$$

The Hölder space $C^{0, \frac{1}{2}}(X; \mathcal{V})$ is defined as

$$\left\{ f \in C(X; \mathcal{V}) \left| \sup_{x, y \in X, x \neq y} \left(\frac{\|f(x) - f(y)\|_{\mathcal{V}}}{\sqrt{d_X(x, y)}} \right) < \infty \right. \right\}.$$

Here, $C^{0,\frac{1}{2}}(X; \mathcal{V})$ is a normed space equipped with the norm

$$\|f\|_{C^{0,\frac{1}{2}}(X; \mathcal{V})} = \|f\|_{C(X; \mathcal{V})} + \sup_{x, y \in X, x \neq y} \left(\frac{\|f(x) - f(y)\|_{\mathcal{V}}}{\sqrt{d_X(x, y)}} \right).$$

Obviously, $C^{0,\frac{1}{2}}(X; \mathcal{V}) \subset C(X; \mathcal{V})$. For simplicity, we denote $C(X) = C(X; \mathbb{R})$ and $C^{0,\frac{1}{2}}(X) = C^{0,\frac{1}{2}}(X; \mathbb{R})$ (see Chapter 5.1 of [Evans \(2010\)](#)). For a given $T > 0$ (see Eq. (3.2.1)), we denote the separable Banach space $C([0, T])$ as $\mathcal{B}$.

- iii) All derivatives implemented in this paper are *weak derivatives* (defined in Chapter 5.2.1 of [Evans \(2010\)](#)). The inner product of $\mathcal{H} = H_0^1([0, T])$, where $\mathcal{H}$ is a separable Hilbert space (see Chapter 8.3 of [Brezis \(2011\)](#)), is defined as

$$\langle f, g \rangle_{\mathcal{H}} = \int_0^T f'(t)g'(t)dt.$$

Unless otherwise stated, the inner product $\langle \cdot, \cdot \rangle$ denotes $\langle \cdot, \cdot \rangle_{\mathcal{H}}$ for simplicity.

- iv) Suppose (X, d_X) is a metric space. $\mathcal{B}(X)$ and $\mathcal{B}(d_X)$ denote the Borel algebra generated by the metric topology corresponding to d_X .

The following inequalities are also useful for deriving many results presented in this paper

$$\|f\|_{\mathcal{B}} \leq \|f\|_{C^{0,\frac{1}{2}}([0, T])} \leq \tilde{C}_T \|f\|_{\mathcal{H}}, \quad \text{for all } f \in \mathcal{H}, \quad (3.1.1)$$

where $\tilde{C}_T$ is a constant depending only on T . The first inequality in Eq. (3.1.1) results from the definition of $\|\cdot\|_{\mathcal{B}}$ and $\|\cdot\|_{C^{0,\frac{1}{2}}([0, T])}$, while the second inequality is from [Evans \(2010\)](#) (Theorem 5 of Chapter 5.6). Eq. (3.1.1) further implies the following embeddings

$$H_0^1([0, T]) = \mathcal{H} \subset C^{0,\frac{1}{2}}([0, T]) \subset \mathcal{B} = C([0, T]). \quad (3.1.2)$$

The algebraic topology concepts referred to in this paper (e.g., Betti numbers, homology groups, and homotopy equivalence) can be found in [Hatcher \(2002\)](#) and [Munkres \(2018\)](#).

We organize this paper as follows. In Section 3.2, we provide a collection of shapes and show that the SECT is well defined for elements in this collection. Additionally, we provide several properties of the SECT using Sobolev and Hölder spaces. These properties will be necessary, in later sections, for developing the probability theory of both the SECT and SECT-based hypothesis testing. In Section 3.3, we construct a probability space for modeling the distributions of shapes. Based on this probability space, we model the SECT as a $C(\mathbb{S}^{d-1}; \mathcal{H})$ -valued random variable. Then, we define the mean and covariance of the SECT. In Section 3.4, we propose the Karhunen–Loève expansion of the SECT based on the mean and covariance defined in Section 3.3, and this expansion leads to

a normal distribution-based statistic for testing hypotheses on the distributions of shapes. In Section 3.5, we provide simulation studies showing the performance of the proposed hypothesis testing approach. In Section 3.6, we conclude this paper and propose several relevant topics for future research. The Appendix in the Supplementary Material (Meng et al., 2022b) provides the necessary mathematical tools for this paper. To avoid distraction from the main flow of the paper, unless otherwise stated, the proof of each theorem is in the Appendix. .

3.2 The Smooth Euler Characteristic Transform of Shapes

In this section, we give the background on the smooth Euler characteristic transform (SECT) and propose corresponding mathematical foundations under the assumption that shapes are deterministic.

3.2.1 The Definition of Smooth Euler Characteristic Transform

We assume $K \subset \overline{B(0, R)} = \{x \in \mathbb{R}^d : \|x\| \leq R\}$ which denotes a closed ball centered at the origin with some radius $R > 0$. For each direction $\nu \in \mathbb{S}^{d-1}$, we define a collection of sublevel sets of K by

$$K_t^\nu \stackrel{\text{def}}{=} \{x \in K | x \cdot \nu \leq t - R\}, \quad \text{for all } t \in [0, T], \quad \text{where } T = 2R. \quad (3.2.1)$$

We then have the following Euler characteristic curve (ECC) in direction ν

$$\chi_t^\nu(K) \stackrel{\text{def}}{=} \text{the Euler characteristic of } K_t^\nu = \chi(K_t^\nu) = \sum_{k=0}^{d-1} (-1)^k \cdot \beta_k(K_t^\nu), \quad (3.2.2)$$

for $t \in [0, T]$, where $\beta_k(K_t^\nu)$ is the k -th Betti number of K_t^ν . The Euler characteristic transform (ECT) defined as $ECT(K) : \mathbb{S}^{d-1} \rightarrow \mathbb{Z}^{[0, T]}, \nu \mapsto \{\chi_t^\nu(K)\}_{t \in [0, T]}$ was first proposed by Turner et al. (2014b) as an alternative to the PHT. Based on the ECT, Crawford et al. (2020) further proposed the SECT as follows

$$\begin{aligned} SECT(K) : \mathbb{S}^{d-1} &\rightarrow \mathbb{R}^{[0, T]}, \quad \nu \mapsto SECT(K)(\nu) = \{SECT(K)(\nu; t)\}_{t \in [0, T]}, \\ \text{where } SECT(K)(\nu; t) &\stackrel{\text{def}}{=} \int_0^t \chi_\tau^\nu(K) d\tau - \frac{t}{T} \int_0^T \chi_\tau^\nu(K) d\tau. \end{aligned} \quad (3.2.3)$$

In order for the integrals in Eq. (3.2.3) to be well-defined, we need to introduce additional conditions. We first propose the following concept motivated by Cohen-Steiner et al. (2007).

Definition 3.2.1. Suppose $K \in \mathcal{S}_d$ and $K \subset \overline{B(0, R)}$. (i) If $K_{t^*}^\nu$ is not homotopy equivalent to $K_{t^*-\delta}^\nu$ for any $\delta > 0$, we call $t^* \in [0, T]$ a homotopy critical point (HCP) of K in direction ν . (ii) If K has finite HCPs in every direction $\nu \in \mathbb{S}^{d-1}$, we call K tame.

Because the Euler characteristic is homotopy invariant, for each $\nu \in \mathbb{S}^{d-1}$, the ECC $\chi_t^\nu(K)$ is a step function of t with finite discontinuities, and hence measurable, if K is tame. These discontinuities are the HCPs of K in

direction ν . We may refine the definition of the ECC $\{\chi_t^\nu(K)\}_{t \in [0, T]}$ in Eq. (3.2.2) as follows: for a tame $K \in \mathcal{S}_d$, define $\chi_t^\nu(K) = \chi(K_t^\nu)$ if t is not an HCP in direction ν , and $\chi_t^\nu(K) = \lim_{t' \rightarrow t+} \chi_{t'}^\nu(K)$ if t is an HCP in direction ν . Under these conditions, $\{\chi_t^\nu(K)\}_{t \in [0, T]}$ is a right continuous with left limits (RCLL) function and has finite discontinuities. We can define k -th Betti number curves $\{\beta_k(K_t^\nu)\}_{t \in [0, T]}$ to be RCLL functions of $t \in [0, T]$ in the same way. To investigate $SECT(K)$ across shapes K , especially the distribution of $SECT(K)$ across K , we need the following condition (needed in the proofs of several theorems in this paper) for tame shapes $K \subset \overline{B(0, R)}$ to restrict our attention to a specific subset of $\mathcal{S}_d$

$$\sup_{k \in \{0, \dots, d-1\}} \left[\sup_{\nu \in \mathbb{S}^{d-1}} \left(\# \left\{ \xi \in \text{Dgm}_k(K; \phi_\nu) \mid \text{pers}(\xi) > 0 \right\} \right) \right] \leq \frac{M}{d}, \quad (3.2.4)$$

where $\text{Dgm}_k(K; \phi_\nu)$ is a PD, $\text{pers}(\xi)$ is the persistence of the homology feature ξ , $\#\{\cdot\}$ denotes the cardinality of a multiset, and $M > 0$ is allowed to be any sufficiently large fixed number. To avoid distraction from the main flow of the paper, we save the details of condition (3.2.4) and the definitions of $\text{Dgm}_k(K; \phi_\nu)$ and $\text{pers}(\xi)$ for the Appendix 3.A. A heuristic interpretation of condition (3.2.4) is that there exists a uniform upper bound for the numbers of nontrivial homology features of shape K across all levels t and directions ν (see Theorem 3.2.1 below). This condition is usually satisfied in medical imaging studies (e.g., a tumor has finitely many connected components and cavities across all levels and directions). Condition (3.2.4) is needed in the proofs of several theorems in this paper, particularly Theorem 3.2.3. Throughout this paper, we focus on shapes in the following collection

$$\mathcal{S}_{R,d}^M \stackrel{\text{def}}{=} \left\{ K \in \mathcal{S}_d \mid K \subset \overline{B(0, R)}, \text{ } K \text{ is tame and satisfies condition (3.2.4) with fixed } M > 0 \right\}.$$

Before Appendix 3.A, we only need the following boundedness derived from condition (3.2.4).

Theorem 3.2.1. *For any $K \in \mathcal{S}_{R,d}^M$, we have the following bounds:*

- (i) $\sup_{\nu \in \mathbb{S}^{d-1}} \left[\sup_{0 \leq t \leq T} \left(\sup_{k \in \{0, \dots, d-1\}} \beta_k(K_t^\nu) \right) \right] \leq M/d$, and
- (ii) $\sup_{\nu \in \mathbb{S}^{d-1}} \left(\sup_{0 \leq t \leq T} |\chi_t^\nu(K)| \right) \leq M$.

The tameness of K and boundedness of $\chi_t^\nu(K)$ in Theorem 3.2.1 guarantee that, for each fixed direction $\nu \in \mathbb{S}^{d-1}$, an ECC $\chi_{(\cdot)}^\nu(K)$ is a measurable and bounded function. Therefore, the integrals in Eq. (3.2.3) are well defined Lebesgue integrals. Since $\{\chi_t^\nu(K)\}_{t \in [0, T]} \in L^1([0, T])$, $SECT(K)(\nu)$ is absolutely continuous on $[0, T]$. Furthermore, we have the following regularity result of the Sobolev type on $SECT(K)(\nu)$ defined in Eq. (3.2.3).

Theorem 3.2.2. *For any $K \in \mathcal{S}_{R,d}^M$ and $\nu \in \mathbb{S}^{d-1}$, we have the following:*

- (i) Function $\{\int_0^t \chi_\tau^\nu(K) d\tau\}_{t \in [0, T]}$ has its first-order weak derivative $\{\chi_t^\nu(K)\}_{t \in [0, T]}$;
- (ii) $SECT(K)(\nu) \in W_0^{1,p}([0, 1]) \subset \mathcal{B}$ for all $p \in [1, \infty)$.

Here, $W_0^{1,p}([0, T])$ is a Sobolev space defined by the following (see the Theorem 8.12 of Brezis (2011) for this specific

definition)

$$W_0^{1,p}([0, T]) = \left\{ f \in L^p([0, T]) \mid \text{weak derivative } f' \text{ exists, } f' \in L^p([0, T]), \text{ and } f(0) = f(T) = 0 \right\}$$

The special case when $p = 2$ in Theorem 3.2.2 is of importance, as $\mathcal{H} = H_0^1([0, T]) = W_0^{1,2}([0, T])$ is the RKHS reproduced by the following kernel

$$\kappa(s, t) = \min\{s, t\} - \frac{st}{T}, \quad \text{for all } s, t \in [0, T], \quad (3.2.5)$$

which is the covariance function of the Brownian bridge on $[0, T]$ (see Example 4.9 of Lifshits (2012)). The relationship between $\mathcal{H}$ and $\kappa(s, t)$ from the RKHS perspective implies that $\langle \kappa(t, \cdot), f \rangle = f(t)$ for all $f \in \mathcal{H}$ and $t \in [0, T]$. This result plays an important role later in this paper. The case when $p = 2$ in Theorem 3.2.2 indicates that $SECT(\mathcal{S}_{R,d}^M) \subset \mathcal{H}^{\mathbb{S}^{d-1}} = \{\text{all maps } F : \mathbb{S}^{d-1} \rightarrow \mathcal{H}\}$, which is enhanced by the following theorem.

Theorem 3.2.3. *For each $K \in \mathcal{S}_{R,d}^M$, we have the following:*

(i) *There exists a constant $C_{M,R,d}^*$ depending only on M , R , and d such that the following two inequalities hold for any two directions $\nu_1, \nu_2 \in \mathbb{S}^{d-1}$,*

$$\begin{aligned} \left(\int_0^T \left| \chi_\tau^{\nu_1}(K) - \chi_\tau^{\nu_2}(K) \right|^2 d\tau \right)^{1/2} &\leq C_{M,R,d}^* \cdot \sqrt{\|\nu_1 - \nu_2\|}, \\ \left\| SECT(K)(\nu_1) - SECT(K)(\nu_2) \right\|_{\mathcal{H}} &\leq C_{M,R,d}^* \cdot \sqrt{\|\nu_1 - \nu_2\| + \|\nu_1 - \nu_2\|^2}. \end{aligned} \quad (3.2.6)$$

(ii) *$SECT(K) \in C^{0, \frac{1}{2}}(\mathbb{S}^{d-1}; \mathcal{H})$, where $\mathbb{S}^{d-1}$ is equipped with the geodesic distance $d_{\mathbb{S}^{d-1}}$.*

(iii) *The constant $\tilde{C}_T$ in Eq. (3.1.1) provides the following inequality*

$$\begin{aligned} &\left| SECT(K)(\nu_1; t_1) - SECT(K)(\nu_2; t_2) \right| \\ &\leq \tilde{C}_T \left\{ \|SECT\|_{C(\mathbb{S}^{d-1}; \mathcal{H})} \cdot \sqrt{|t_1 - t_2|} + C_{M,R,d}^* \cdot \sqrt{\|\nu_1 - \nu_2\| + \|\nu_1 - \nu_2\|^2} \right\}, \end{aligned} \quad (3.2.7)$$

for all $\nu_1, \nu_2 \in \mathbb{S}^{d-1}$ and $t_1, t_2 \in [0, T]$, which implies that $(\nu, t) \mapsto SECT(K)(\nu; t)$, as a function on $\mathbb{S}^{d-1} \times [0, T]$, belongs to $C^{0, \frac{1}{2}}(\mathbb{S}^{d-1} \times [0, T]; \mathbb{R})$.

Theorem 3.2.3 (i) is a counterpart of the Lemma 2.1 in Turner et al. (2014b) and is derived using “bottleneck stability” (see Cohen-Steiner et al. (2007)). Additionally, Theorem 3.2.3 (i) is implemented in the proof of Theorem 3.2.3 (ii) and will be used in Section 3.2.3. Furthermore, Theorem 3.2.3 (ii) guarantees that $SECT(\mathcal{S}_{R,d}^M) \subset C^{0, \frac{1}{2}}(\mathbb{S}^{d-1}; \mathcal{H}) \subset C(\mathbb{S}^{d-1}; \mathcal{H}) \subset \mathcal{H}^{\mathbb{S}^{d-1}}$. As a result, Eq. (3.2.3) defines the following map

$$SECT : \mathcal{S}_{R,d}^M \rightarrow C(\mathbb{S}^{d-1}; \mathcal{H}), \quad K \mapsto \{SECT(K)(\nu)\}_{\nu \in \mathbb{S}^{d-1}} \stackrel{\text{def}}{=} SECT(K). \quad (3.2.8)$$

Theorem 3.2.3 (iii) will be implemented in the mathematical foundation of the hypothesis testing in Section 3.4 (specifically, see Theorems 3.3.4 and 3.4.1).

Corollary 6 of Ghrist et al. (2018) implies the following result, which shows that the map in Eq. (3.2.8) preserves all the information of shapes $K \in \mathcal{S}_{R,d}^M$.

Theorem 3.2.4. *The map SECT defined in Eq. (3.2.8) is invertible for all dimensions d .*

Together with Eq. (3.2.8), Theorem 3.2.4 enables us to view each shape $K \in \mathcal{S}_{R,d}^M$ as $SECT(K)$ which is an element of $C(\mathbb{S}^{d-1}; \mathcal{H})$. This viewpoint will help us characterize the randomness of shapes $K \in \mathcal{S}_{R,d}^M$ using probability measures on the separable Banach space $C(\mathbb{S}^{d-1}; \mathcal{H})$.

If the shapes of interest are represented as triangular meshes, their ECC defined in Eq. (3.2.2) can easily be computed (e.g., Section 2 of Wang et al. (2021)). The SECT is then derived directly from the computed ECC (see Eq. (3.2.3)). In Appendix 3.B, we briefly introduce an approach for approximating the ECC of shapes in $\mathbb{R}^d$ in scenarios where the triangular mesh representation of shapes is not available, which in turn provides a corresponding approximation to SECT. This approach applies Čech complexes and the nerve theorem (see Chapter III of Edelsbrunner and Harer (2010)), which are implemented in all proof-of-concept examples and simulations throughout this paper. In addition, an approach was studied in Niyogi et al. (2008) for finding the Betti numbers of submanifolds of Euclidean spaces from random point clouds. For the case where a point cloud is drawn near a submanifold, this approach can be applied to estimate the SECT of the submanifold. The application of Niyogi et al. (2008) to SECT is left for future research.

3.2.2 Proof-of-Concept Simulation Examples I: Deterministic Shapes

In this subsection, we compute the SECT of two simulated shapes $K^{(1)}$ and $K^{(2)}$ of dimension $d = 2$. These shapes are defined as the following and presented in Figures 3.1(a) and (g), respectively:

$$K^{(j)} = \left\{ x \in \mathbb{R}^2 \left| \inf_{y \in S^{(j)}} \|x - y\| \leq \frac{1}{5} \right. \right\}, \quad \text{where } j \in \{1, 2\}, \quad (3.2.9)$$

$$\begin{aligned} S^{(1)} &= \left\{ \left(\frac{2}{5} + \cos t, \sin t \right) \left| \frac{\pi}{5} \leq t \leq \frac{9\pi}{5} \right. \right\} \cup \left\{ \left(-\frac{2}{5} + \cos t, \sin t \right) \left| \frac{6\pi}{5} \leq t \leq \frac{14\pi}{5} \right. \right\}, \\ S^{(2)} &= \left\{ \left(\frac{2}{5} + \cos t, \sin t \right) \left| 0 \leq t \leq 2\pi \right. \right\} \cup \left\{ \left(-\frac{2}{5} + \cos t, \sin t \right) \left| \frac{6\pi}{5} \leq t \leq \frac{14\pi}{5} \right. \right\}. \end{aligned}$$

We compute $SECT(K^{(j)})(\nu; t)$ across direction $\nu \in \mathbb{S}^1$ and sublevel set $t \in [0, 3]$ with $T = 3$. For the following visualization, we identify each direction $\nu \in \mathbb{S}^1$ through the parametrization $\nu = (\cos \vartheta, \sin \vartheta)$ with $\vartheta \in [0, 2\pi)$.

- The surfaces of the bivariate maps $(\vartheta, t) \mapsto SECT(K^{(j)})(\nu; t)$, for $j \in \{1, 2\}$, are presented in Figures 3.1(b), (c), (h), and (i).

- The curves of the univariate maps $t \mapsto SECT(K^{(j)})((1,0)\mathbf{I}; t)$, for $j \in \{1, 2\}$, are presented by the black solid lines in Figures 3.1(d) and (j); while the curves of the univariate maps $t \mapsto SECT(K^{(j)})((0,1)\mathbf{I}; t)$, for $j \in \{1, 2\}$, are presented by the black solid lines in Figures 3.1(e) and (k).
- Lastly, the curves of the univariate maps $\vartheta \mapsto SECT(K^{(j)})((\cos \vartheta, \sin \vartheta)\mathbf{I}; \frac{3}{2})$, for $j \in \{1, 2\}$, are presented by the black solid lines in Figures 3.1(f) and (l).

These figures illustrate the continuity of $(\nu, t) \mapsto SECT(K^{(j)})(\nu; t)$ stated in Theorem 3.2.3 (iii). Specifically, the curves and surfaces in these figures look not as rough as the path of Brownian motions, while they are not differentiable everywhere. With probability one, paths of a Brownian motion are not locally $C^{0, \frac{1}{2}}$ -continuous (see Remark 22.4 of Klenke (2013)). Hence, based on Figure 3.1, the regularity of $(\nu, t) \mapsto SECT(K^{(j)})(\nu; t)$ is likely to be better than that of Brownian motion paths, but worse than continuously differentiable functions. Therefore, Figure 3.1 supports the $C^{0, \frac{1}{2}}$ -continuity in Theorem 3.1.

The invertibility in Theorem 3.2.4 indicates that all information of $K^{(1)}$ and $K^{(2)}$ is stored in the surfaces presented by Figures 3.1(b), (c), (h), and (i). The red dashed curves in Figures 3.1(j), (k), and (l) are the counterparts of $K^{(1)}$ (see the curves in Figures 3.1(d), (e), and (f)). The discrepancy between the solid black and dashed red curves illustrates the SECT's ability to distinguish shapes. Simulation studies in Section 3.5 will further confirm the ability of SECT in distinguishing shapes.

3.2.3 The Primitive Euler Characteristic Transform of Shapes

In this work, we also introduce another topological invariant — the primitive Euler characteristic transform (PECT) — which is related to SECT and will play an important role in defining a distance function on $\mathcal{S}_{R,d}^M$. The PECT is defined as follows

$$PECT : \mathcal{S}_{R,d}^M \rightarrow C(\mathbb{S}^{d-1}; \mathcal{H}_{BM}), \quad K \mapsto PECT(K) \stackrel{\text{def}}{=} \{PECT(K)(\nu)\}_{\nu \in \mathbb{S}^{d-1}},$$

$$\text{where } PECT(K)(\nu) \stackrel{\text{def}}{=} \left\{ \int_0^t \chi_\tau^\nu(K) d\tau \right\}_{t \in [0, T]}, \quad (3.2.10)$$

and $\mathcal{H}_{BM} \stackrel{\text{def}}{=} \{f \in L^2([0, T]) \mid \text{weak derivative } f' \text{ exists, } f' \in L^2([0, T]), \text{ and } x(0) = 0\}$ is a separable Hilbert space equipped with the inner product $\langle f, g \rangle_{\mathcal{H}_{BM}} = \int_0^T f'(t)g'(t)dt$. In addition, $\mathcal{H}_{BM}$ is the RKHS generated by the covariance function $(s, t) \mapsto \min\{s, t\}$ of the Brownian motion on $[0, T]$. The inequality in Eq. (3.2.6), together with the weak derivative of $PECT(K)(\nu)$ being $\{\chi_t^\nu(K)\}_{t \in [0, T]}$ (see Theorem 3.2.2), implies that

$$\|PECT(K)(\nu_1) - PECT(K)(\nu_2)\|_{\mathcal{H}_{BM}} \leq C_{M,R,d}^* \cdot \sqrt{\|\nu_1 - \nu_2\|}.$$

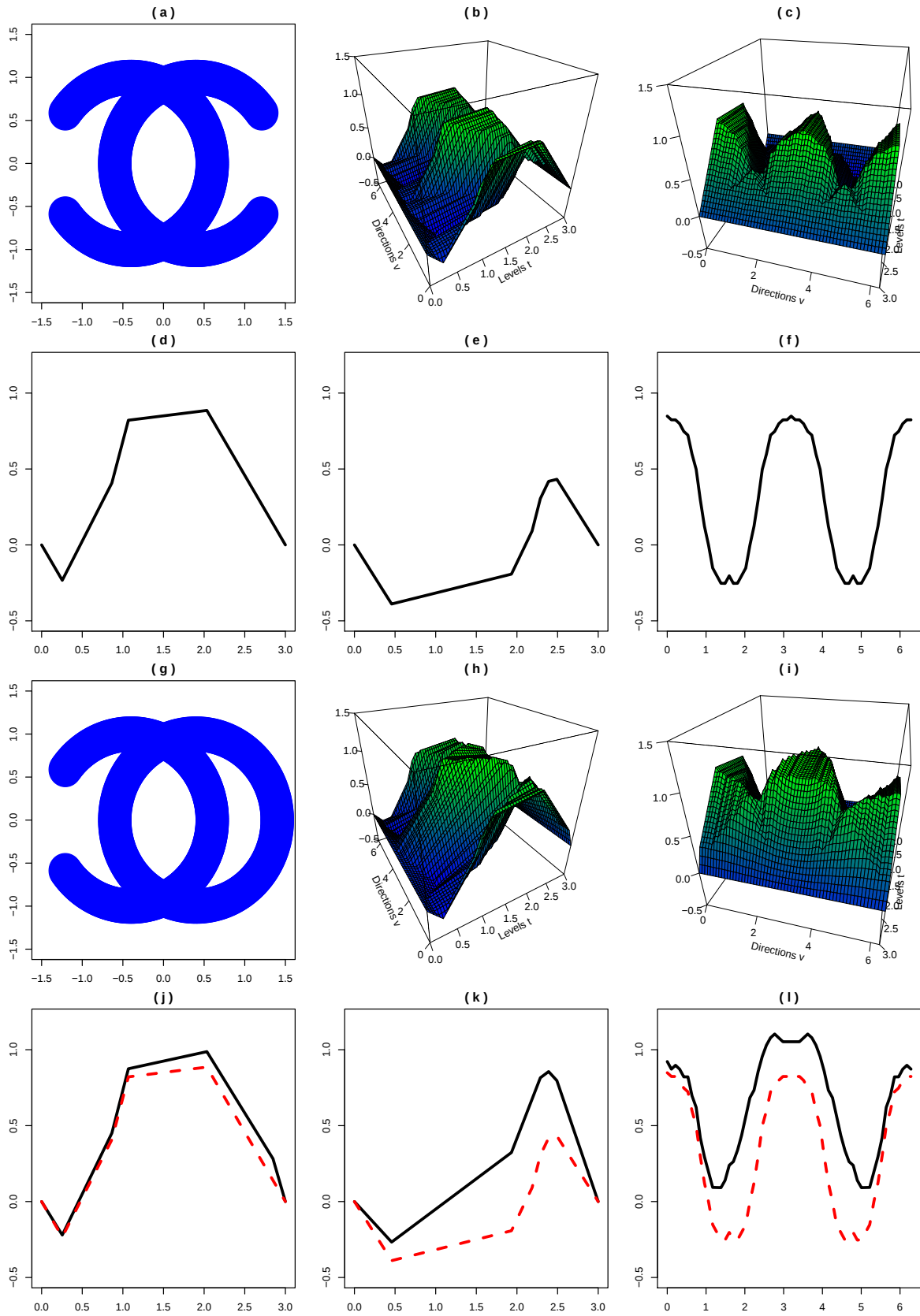


Figure 3.1: Visualizations of $SECT(K^{(j)})$ for $j \in \{1, 2\}$, where $K^{(1)}$ and $K^{(2)}$ are defined in Eq. (3.2.9). Panels (b) and (c) present the same surface from different angles. Similarly, panels (h) and (i) present the same surface. The similarity between the curves in the panel (j) partially indicates that SECT in only one direction does not preserve all the geometric information of a shape.

Therefore, we have the $PECT(K) \in C(\mathbb{S}^{d-1}; \mathcal{H}_{BM})$ in Eq. (3.2.10). $PECT$ relates to $SECT$ via the following equation.

$$SECT(K)(\nu; t) = PECT(K)(\nu; t) - \frac{t}{T} PECT(K)(\nu; t), \quad (3.2.11)$$

for all $\nu \in \mathbb{S}^{d-1}$ and $t \in [0, T]$. Additionally, Theorem 3.2.4, together with Eq. (3.2.11), implies that the $PECT$ in Eq. (3.2.10) is invertible.

3.3 Random Shapes and Probabilistic Distributions over the SECT

In Section 3.2, shapes K were viewed as deterministic. From this section onward, we investigate the randomness of shapes. Suppose now that $\mathcal{S}_{R,d}^M$ is equipped with a σ -algebra $\mathcal{F}$ and that a distribution of shapes K across $\mathcal{S}_{R,d}^M$ can be represented by a probability measure $\mathbb{P} = \mathbb{P}(dK)$ on $\mathcal{F}$. Then, $(\mathcal{S}_{R,d}^M, \mathcal{F}, \mathbb{P})$ is a probability space, and it is the underlying probability space of interest.

3.3.1 Formulation of the SECT as a Random Variable

For each direction $\nu \in \mathbb{S}^{d-1}$ and level $t \in [0, T]$, the integer-valued map $\chi_t^\nu : K \mapsto \chi_t^\nu(K)$ is defined on $\mathcal{S}_{R,d}^M$. We need the following assumption on the measurability of χ_t^ν .

Assumption 1. *For each fixed $\nu \in \mathbb{S}^{d-1}$ and $t \in [0, T]$, the map*

$$\chi_t^\nu : (\mathcal{S}_{R,d}^M, \mathcal{F}) \rightarrow (\mathbb{R}, \mathcal{B}(\mathbb{R}))$$

is a real-valued random variable — meaning that $(\chi_t^\nu)^{-1}(B) = \{K \in \mathcal{S}_{R,d}^M \mid \chi_t^\nu(K) \in B\} \in \mathcal{F}$ for all $B \in \mathcal{B}(\mathbb{R})$.

A σ -algebra $\mathcal{F}$ satisfying Assumption 1 does exist. Specifically, we construct a metric on $\mathcal{S}_{R,d}^M$ and show that the Borel algebra induced by this metric satisfies Assumption 1. Motivated by the metric of $C(\mathbb{S}^{d-1}; \mathcal{H}_{BM})$, we define the semi-distance

$$\begin{aligned} \rho : \mathcal{S}_{R,d}^M \times \mathcal{S}_{R,d}^M &\rightarrow [0, \infty), \quad (K_1, K_2) \mapsto \rho(K_1, K_2), \quad \text{where} \\ \rho(K_1, K_2) &\stackrel{\text{def}}{=} \sup_{\nu \in \mathbb{S}^{d-1}} \left\{ \left(\int_0^T \left| \chi_\tau^\nu(K_1) - \chi_\tau^\nu(K_2) \right|^2 d\tau \right)^{1/2} \right\}. \end{aligned} \quad (3.3.1)$$

The following theorem shows that the ρ defined in Eq. (3.3.1) is indeed a metric on $\mathcal{S}_{R,d}^M$. Furthermore, it provides an example of a σ -algebra satisfying Assumption 1.

Theorem 3.3.1. *(i) The map ρ defined in Eq. (3.3.1) is a metric on $\mathcal{S}_{R,d}^M$.*

(ii) Assumption 1 is satisfied if $\mathcal{F} = \mathcal{B}(\rho)$.

Additionally, this ρ metric is a counterpart of the $dist_{\mathcal{M}_d}$ metric in [Turner et al. \(2014b\)](#) (Eq. (2.3) therein) and the $dist_{\mathcal{M}_{d-1}}^{SECT}$ metric in [Crawford et al. \(2020\)](#) (Eq. (7) therein).

Under Assumption 1, the ECC $\{\chi_t^\nu\}_{t \in [0, T]}$, for each fixed $\nu \in \mathbb{S}^{d-1}$, is a continuous-time stochastic process with RCLL paths defined on probability space $(\mathcal{S}_{R,d}^M, \mathcal{F}, \mathbb{P})$. Since each path $\{\chi_t^\nu(K)\}_{t \in [0, T]}$ is a step function with finitely many discontinuities, the integrals $\int_0^t \chi_\tau^\nu(K) d\tau$ for $t \in [0, T]$ are Riemann integrals of the form

$$\int_0^t \chi_\tau^\nu(K) d\tau = \lim_{n \rightarrow \infty} \left\{ \frac{t}{n} \sum_{l=1}^n \chi_{\frac{tl}{n}}^\nu(K) \right\}, \quad \text{for all } K \in \mathcal{S}_{R,d}^M. \quad (3.3.2)$$

Because each $\chi_{\frac{tl}{n}}^\nu$ in Eq. (3.3.2) is a random variable under Assumption 1, the limit in Eq. (3.3.2) for each fixed $t \in [0, T]$ is a random variable as well. Therefore, for each fixed $\nu \in \mathbb{S}^{d-1}$, $\{\int_0^t \chi_\tau^\nu d\tau\}_{t \in [0, T]}$ with $\int_0^t \chi_\tau^\nu d\tau : K \mapsto \int_0^t \chi_\tau^\nu(K) d\tau$ is a real-valued stochastic process with continuous paths. Then, under Assumption 1, Eqs. (3.2.3) and (3.2.10) define the following real-valued stochastic processes on the probability space $(\mathcal{S}_{R,d}^M, \mathcal{F}, \mathbb{P})$

$$\begin{aligned} SECT(\cdot)(\nu) &\stackrel{\text{def}}{=} \left\{ \int_0^t \chi_\tau^\nu d\tau - \frac{t}{T} \int_0^T \chi_\tau^\nu d\tau \stackrel{\text{def}}{=} SECT(\cdot)(\nu; t) \right\}_{t \in [0, T]}, \\ PECT(\cdot)(\nu) &\stackrel{\text{def}}{=} \left\{ \int_0^t \chi_\tau^\nu d\tau \stackrel{\text{def}}{=} PECT(\cdot)(\nu; t) \right\}_{t \in [0, T]}. \end{aligned} \quad (3.3.3)$$

Eq. (3.2.11) indicates that $SECT(\cdot)(\nu)$ is the “bridge version”[‡] of $PECT(\cdot)(\nu)$. Theorem 3.2.2, together with [Berlinet and Thomas-Agnan \(2011\)](#) (see Corollary 13 in Chapter 4 therein), directly implies the following theorem; hence, the corresponding proof is omitted.

Theorem 3.3.2. *For each direction $\nu \in \mathbb{S}^{d-1}$, under Assumption 1, $SECT(\cdot)(\nu)$ is a real-valued stochastic process with paths in $\mathcal{H}$. Equivalently, $SECT(\cdot)(\nu)$ is a random variable taking values in $(\mathcal{H}, \mathcal{B}(\mathcal{H}))$. Additionally, $SECT(\cdot)(\nu; 0) = SECT(\cdot)(\nu; T) = 0$.*

3.3.2 Mean and Covariance of the SECT

Theorem 3.3.2 indicates that the SECT is an $\mathcal{H}$ -valued random field on manifold $\mathbb{S}^{d-1}$. Theorem 3.2.3 implies that the sample paths of this random field are in $C^{0, \frac{1}{2}}(\mathbb{S}^{d-1}; \mathcal{H})$. To define the mean and covariance for the SECT, we need the following assumption on the second moments corresponding to the distribution $\mathbb{P}$.

Assumption 2. *We have the following finite second moments for all $\nu \in \mathbb{S}^{d-1}$*

$$\mathbb{E} \|SECT(\cdot)(\nu)\|_{\mathcal{H}}^2 = \int_{\mathcal{S}_{R,d}^M} \|SECT(K)(\nu)\|_{\mathcal{H}}^2 \mathbb{P}(dK) < \infty.$$

[‡]: Motivated by Brownian bridges, for any stochastic process $\{X_t\}_{t \in [0, T]}$ with continuous paths and $X_0 = 0$, we refer to $\{U_t = X_t - \frac{t}{T} X_T\}_{t \in [0, T]}$ as its bridge version.

Together, Eq. (3.1.1) and Assumption 2 imply that

$$\mathbb{E}|SECT(\cdot)(\nu; t)|^2 \leq \tilde{C}_T^2 \cdot \mathbb{E}\|SECT(\cdot)(\nu)\|_{\mathcal{H}}^2 < \infty, \quad \text{for all } \nu \text{ and } t. \quad (3.3.4)$$

Therefore, we may define the mean function $\{m\}_{\nu \in \mathbb{S}^{d-1}}$ as follows under Assumption 2

$$\begin{aligned} m_\nu(t) &= \mathbb{E}\{SECT(\cdot)(\nu; t)\} \\ &= \int_{\mathcal{S}_{R,d}^M} \left\{ \int_0^t \chi_\tau^\nu(K) d\tau - \frac{t}{T} \int_0^T \chi_\tau^\nu(K) d\tau \right\} \mathbb{P}(dK), \quad t \in [0, T]. \end{aligned} \quad (3.3.5)$$

The next theorem is crucial for defining the covariance of the SECT as well as for conducting the hypothesis testing framework that we will detail in Section 3.4.

Theorem 3.3.3. (i) For each fixed direction $\nu \in \mathbb{S}^{d-1}$, the real-valued function $m_\nu \stackrel{\text{def}}{=} \{m_\nu(t)\}_{t \in [0, T]}$ of t belongs to $\mathcal{H}$; hence, $SECT(K)(\nu) - m_\nu \in \mathcal{H}$ for all $K \in \mathcal{S}_{R,d}^M$.

(ii) The map $(K, t) \mapsto SECT(K)(\nu; t)$ belongs to $L^2(\mathcal{S}_{R,d}^M \times [0, T], \mathbb{P}(dK) \times dt)$, where $\mathbb{P}(dK) \times dt$ denotes the product measure generated by $\mathbb{P}(dK)$ and the Lebesgue measure dt .

(iii) The map $\nu \mapsto m_\nu$ belongs to $C^{0, \frac{1}{2}}(\mathbb{S}^{d-1}; \mathcal{H})$; hence, this map belongs to $C(\mathbb{S}^{d-1}; \mathcal{H})$.

As an analog to the covariance of Gaussian measures on separable Banach spaces (see Hairer (2009), Eq. (3.1) therein), the covariance of the SECT is defined as a linear map $C(\nu_1, \nu_2) : \mathcal{H} \rightarrow \mathcal{H}$ via the following equation

$$\begin{aligned} &\langle h_1, C(\nu_1, \nu_2) h_2 \rangle \\ &\stackrel{\text{def}}{=} \mathbb{E} \left\{ \langle h_1, SECT(\cdot)(\nu_1) - m_{\nu_1} \rangle \cdot \langle h_1, SECT(\cdot)(\nu_2) - m_{\nu_2} \rangle \right\}, \end{aligned} \quad (3.3.6)$$

for all $h_1, h_2 \in \mathcal{H}$. Eq. (3.3.6) is well-defined because of the first statement in Theorem 3.3.3 and the following inequalities induced by Assumption 2

$$\begin{aligned} &\mathbb{E} \left| \langle h_1, SECT(\cdot)(\nu_1) - m_{\nu_1} \rangle \cdot \langle h_1, SECT(\cdot)(\nu_2) - m_{\nu_2} \rangle \right| \\ &\leq \|h_1\|_{\mathcal{H}} \|h_2\|_{\mathcal{H}} \cdot \left(\mathbb{E}\|SECT(\cdot)(\nu_1) - m_{\nu_1}\|_{\mathcal{H}}^2 \right)^{1/2} \left(\mathbb{E}\|SECT(\cdot)(\nu_2) - m_{\nu_2}\|_{\mathcal{H}}^2 \right)^{1/2} < \infty, \end{aligned}$$

which further indicate that $C(\nu_1, \nu_2)$ is a bounded linear operator.

Because of the finite second moments in Eq. (3.3.4), the covariance function

$$\Xi_\nu(s, t) = \text{cov}(SECT(\cdot)(\nu; s), SECT(\cdot)(\nu; t)), \quad \text{for } s, t \in [0, T],$$

is well-defined. The real-valued covariance $\Xi_\nu(s, t)$ is determined by the operator-valued covariance $C(\nu, \nu)$ via the

following

$$\begin{aligned}\Xi_\nu(s, t) &= \mathbb{E} \left\{ \left\langle \kappa(s, \cdot), SECT(\cdot)(\nu) - m_\nu \right\rangle \cdot \left\langle \kappa(t, \cdot), SECT(\cdot)(\nu) - m_\nu \right\rangle \right\} \\ &= \left\langle \kappa(s, \cdot), C(\nu, \nu) \kappa(t, \cdot) \right\rangle, \quad \text{for all } \nu \in \mathbb{S}^{d-1},\end{aligned}\tag{3.3.7}$$

where the first identity follows from the fact that κ as defined in Eq. (3.2.5) is the reproducing kernel of RKHS $\mathcal{H}$. The following theorem on the covariance function $\Xi_\nu(s, t)$ will be implemented in Section 3.4.

Theorem 3.3.4. *For each fixed direction $\nu \in \mathbb{S}^{d-1}$, we have the following: (i) $SECT(\cdot)(\nu)$ is mean-square continuous (i.e., $\lim_{\epsilon \rightarrow 0} \mathbb{E} |SECT(\cdot)(\nu; t + \epsilon) - SECT(\cdot)(\nu; t)|^2 = 0$); (ii) $(s, t) \mapsto \Xi_\nu(s, t)$ is continuous on $[0, T]^2$.*

The first part of Theorem 3.3.4 directly follows from Eq. (3.2.7), while the second half follows from Lemma 4.2 of Alexanderian (2015); hence, the proof of Theorem 3.3.4 is omitted.

In applications, we cannot reasonably sample infinitely many directions $\nu \in \mathbb{S}^{d-1}$ and levels $t \in [0, T]$. For each given shape K , we can only compute $SECT(K)(\nu; t)$ for finitely many directions $\{\nu_1, \dots, \nu_\Gamma\} \subset \mathbb{S}^{d-1}$ and levels $\{t_1, \dots, t_\Delta\} \subset [0, T]$ with $t_1 < t_2 < \dots < t_\Delta$. Hence, for a collection of shapes $\{K_i\}_{i=1}^n \subset \mathcal{S}_{R,d}^M$ sampled from $\mathbb{P}$, the data we have in applications are presented by the following 3-dimensional arrays

$$\left\{ SECT(K_i)(\nu_p, t_q) \mid i = 1, \dots, n, p = 1, \dots, \Gamma, \text{ and } q = 1, \dots, \Delta \right\}.\tag{3.3.8}$$

To preserve as much information about a shape K as possible, we need to make the numbers of directions and levels (i.e., Γ and Δ in Eq. (3.3.8)) sufficiently large. Curry et al. (2018) (Theorem 7.14 therein) provided an explicit formula of the upper bound on the minimal number of directions. Levels $\{t_j\}_{j=1}^\Delta$ should be sufficiently dense such that there is at least one t_j between any two consecutive HCPs in each direction.

3.3.3 Proof-of-Concept Simulation Examples II: Random Shapes

In this subsection, we compute the SECT for a collection of shapes $\{K_i\}_{i=1}^n$ of dimension $d = 2$. These shapes are randomly generated as follows

$$\begin{aligned}K_i &= \left\{ x \in \mathbb{R}^2 \mid \inf_{y \in S_i} \|x - y\| \leq \frac{1}{5} \right\}, \quad \text{where} \\ S_i &= \left\{ \left(\frac{2}{5} + a_{1,i} \times \cos t, b_{1,i} \times \sin t \right) \mid \frac{\pi}{5} \leq t \leq \frac{9\pi}{5} \right\} \\ &\quad \cup \left\{ \left(-\frac{2}{5} + a_{2,i} \times \cos t, b_{2,i} \times \sin t \right) \mid \frac{6\pi}{5} \leq t \leq \frac{14\pi}{5} \right\},\end{aligned}\tag{3.3.9}$$

and $\{a_{1,i}, a_{2,i}, b_{1,i}, b_{2,i}\}_{i=1}^n \stackrel{i.i.d.}{\sim} N(1, 0.05^2)$. One element of the shape collection $\{K_i\}_{i=1}^n$ is presented in Figure 3.2(a). The underlying distribution on $\mathcal{S}_{R,d}^M$ generating $\{K_i\}_{i=1}^n$ is denoted by $\mathbb{P}$, and the expectation associated with $\mathbb{P}$ is denoted by $\mathbb{E}$. We estimate the expected value $\mathbb{E}\{SECT(\cdot)(\nu; t)\}$ by the sample average $\frac{1}{n} \sum_{i=1}^n SECT(K_i)(\nu; t)$ with $n = 100$. We identify each direction $\nu \in \mathbb{S}^1$ through the parametrization $\nu = (\cos \vartheta, \sin \vartheta)$ with some $\vartheta \in [0, 2\pi)$ as we did in Section 3.2.2.

- The surface of the map $(\vartheta, t) \mapsto \mathbb{E}\{SECT(\cdot)(\nu; t)\}$ is presented in Figures 3.2(b) and (c).
- The black solid curves in Figure 3.2(d) present the 100 paths $t \mapsto SECT(K_i)((1, 0)^\top; t)$, the black solid curves in Figure 3.2(e) present paths $t \mapsto SECT(K_i)((0, 1)^\top; t)$, and the black solid curves in Figure 3.2(f) present paths $\vartheta \mapsto SECT(K_i)((\cos \vartheta, \sin \vartheta)^\top; \frac{3}{2})$, for $i \in \{1, \dots, 100\}$.
- The red solid curves in Figures 3.2(d), (e), and (f) present mean curves $t \mapsto \mathbb{E}\{SECT(\cdot)((1, 0)^\top; t)\}$, $t \mapsto \mathbb{E}\{SECT(\cdot)((0, 1)^\top; t)\}$, and $\vartheta \mapsto \mathbb{E}\{SECT(\cdot)((\cos \vartheta, \sin \vartheta)^\top; \frac{3}{2})\}$, respectively.

The smoothness of the red solid curves in Figures 3.2(d) and (e) supports the regularity of $\{m_\nu(t)\}_{t \in [0, T]}$ in Theorem 3.3.3. The finite variance of $SECT(\cdot)(\nu; t)$ for $\nu = (1, 0)^\top, (0, 1)^\top$ and $t = 3/2$ visually presented in Figures 3.2(d), (e), and (f) supports Eq. (3.3.4).

In addition, the blue dashed curves in Figures 3.2(d), (e), and (f) present curves $t \mapsto SECT(K^{(1)}((1, 0)^\top; t)$, $t \mapsto SECT(K^{(1)}((0, 1)^\top; t)$, and $\vartheta \mapsto SECT(K^{(1)}((\cos \vartheta, \sin \vartheta)^\top; \frac{3}{2}))$, respectively, where shape $K^{(1)}$ is defined in Eq. (3.2.9). Since $\mathbb{E}\{a_{1,i}\} = \mathbb{E}\{a_{2,i}\} = \mathbb{E}\{b_{1,i}\} = \mathbb{E}\{b_{2,i}\} = 1$, the shape $K^{(1)}$ defined in Eq. (3.2.9) can be viewed as the “mean shape” of the random collection $\{K_i\}_{i=1}^n$. The similarity between the red solid curves and blue dashed curves in 3.2(d), (e), and (f) supports the “mean shape” role of $K^{(1)}$. The rigorous definition of a “mean shape” and its relationship to the mean function $\mathbb{E}\{SECT(\cdot)(\nu; t)\}$ is left for future research. A potential way of defining mean shapes is through the following Fréchet mean form (Fréchet, 1948)

$$K_\oplus \stackrel{\text{def}}{=} \arg \min_{K \in \mathcal{S}_{R,d}^M} \mathbb{E} \left[\{\rho(\cdot, K)\}^2 \right] = \arg \min_{K \in \mathcal{S}_{R,d}^M} \left[\int_{\mathcal{S}_{R,d}^M} \{\rho(K', K)\}^2 \mathbb{P}(dK') \right], \quad (3.3.10)$$

where ρ can be either the metric on $\mathcal{S}_{R,d}^M$ defined in Eq. (3.3.1) or any other metrics generating σ -algebras satisfying Assumption 1. The existence and uniqueness of the minimizer $K_\oplus$ in Eq. (3.3.10), the relationship between $SECT(K_\oplus)$ and $\mathbb{E}\{SECT(\cdot)\}$, and the extension of Eq. (3.3.10) to Fréchet regression (Petersen and Müller, 2019) for random shapes are left for future research. The study of the existence of $K_\oplus$ will be a counterpart of Section 4 of Mileyko et al. (2011).

In the scenarios where the SECT of shapes from distribution $\mathbb{P}$ are computed only in finitely many directions and at finitely many levels (see the end of Section 3.3.2), the mean surface $(\vartheta, t) \mapsto \mathbb{E}\{SECT(\cdot)(\nu; t)\}$ in Figures 3.2(b) and (c) can also be potentially estimated using manifold learning methods (e.g., Yue et al. (2016), Dunson and Wu (2021), and Meng and Eloyan (2021b)).

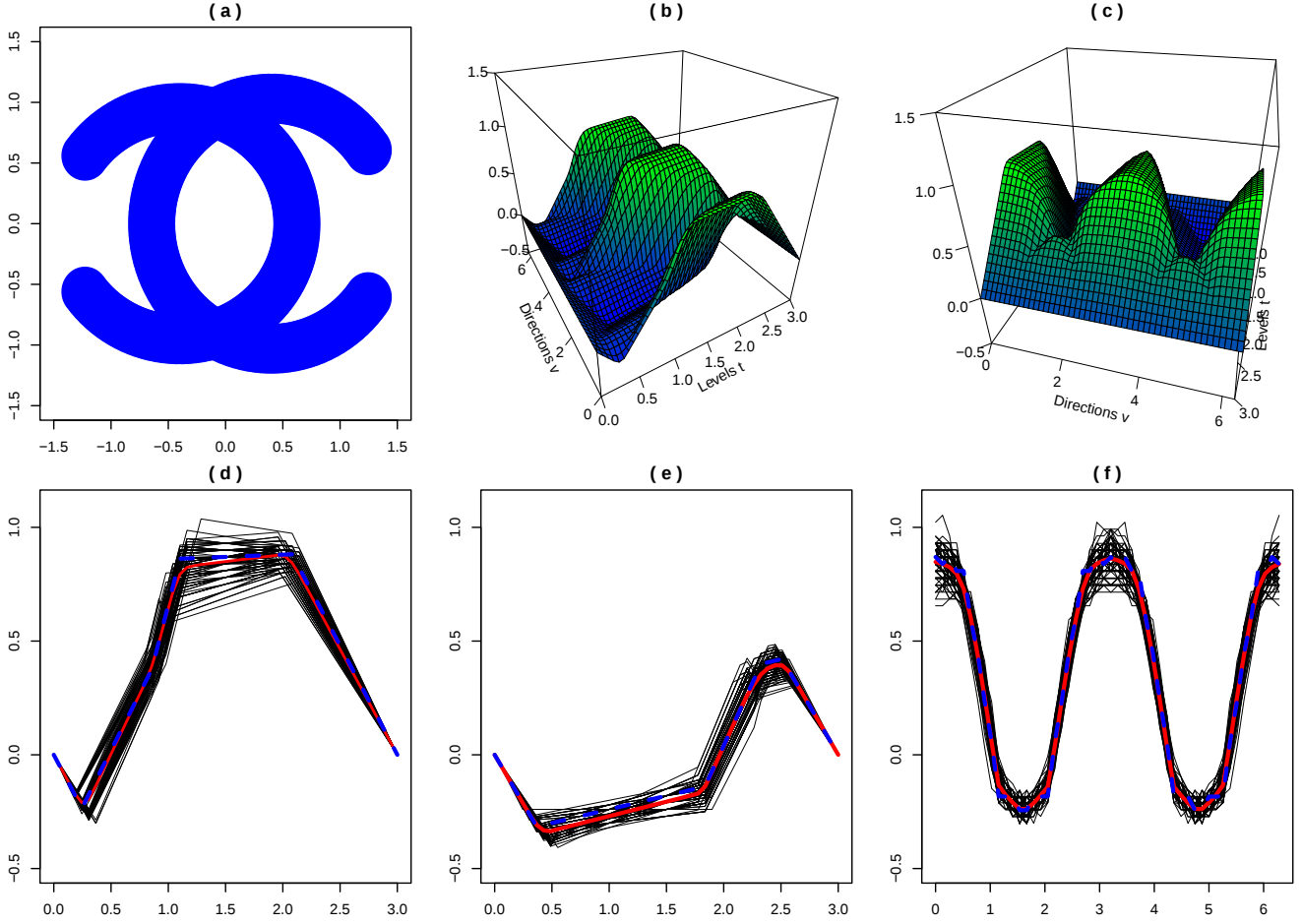


Figure 3.2: Visualizations of $\mathbb{E}\{SECT(\cdot)(\nu; t)\}$, $\{SECT(K_i)\}_{i=1}^n$ with $n = 100$, and $SECT(K^{(1)})(\nu; t)$. Panels (b) and (c) present the same surface, but from different angles.

3.4 Testing Hypotheses on Shapes

In this section, we apply the probabilistic formulation proposed in Section 3.3 to testing statistical hypotheses on shapes. Suppose $\mathbb{P}^{(1)}$ and $\mathbb{P}^{(2)}$ are two independent distributions on $\mathcal{S}_{R,d}^M$ satisfying Assumption 2, and $m_\nu^{(j)}(t) = \int_{\mathcal{S}_{R,d}^M} SECT(K)(\nu; t) \mathbb{P}^{(j)}(dK)$, for $j \in \{1, 2\}$, are the mean functions corresponding to the two distributions. We are interested in testing the following hypotheses

$$\begin{aligned} H_0 : m_\nu^{(1)}(t) &= m_\nu^{(2)}(t) \text{ for all } (\nu, t) \in \mathbb{S}^{d-1} \times [0, T] \\ \text{vs. } H_1 : m_\nu^{(1)}(t) &\neq m_\nu^{(2)}(t) \text{ for some } (\nu, t). \end{aligned} \tag{3.4.1}$$

For example, suppose we are interested in the distributions of heel bones of two groups of primates, testing the hypotheses above helps us distinguish the two groups of primates from a morphology perspective.

The null hypothesis H_0 in Eq. (3.4.1) is equivalent to $\sup_{\nu \in \mathbb{S}^{d-1}} \{ \|m_\nu^{(1)} - m_\nu^{(2)}\|_{\mathcal{B}} \} = 0$. Theorem 3.3.3 (iii),

together with Eq. (3.1.1), indicates that the following maximizer exists

$$\nu^* \stackrel{\text{def}}{=} \arg \max_{\nu \in \mathbb{S}^{d-1}} \left\{ \|m_{\nu}^{(1)} - m_{\nu}^{(2)}\|_{\mathcal{B}} \right\}. \quad (3.4.2)$$

Therefore, H_0 is equivalent to assessing whether $\|m_{\nu^*}^{(1)} - m_{\nu^*}^{(2)}\|_{\mathcal{B}} = 0$. The direction ν^* defined in Eq. (3.4.2) is called a *distinguishing direction*, which we will focus on with the SECT. Hence, we investigate the distributions of the real-valued stochastic process $SECT(\cdot)(\nu^*)$ corresponding to $\mathbb{P}^{(1)}$ and $\mathbb{P}^{(2)}$, respectively.

3.4.1 Karhunen–Loève Expansion

Let $C^{(j)}(\nu_1, \nu_2)$ be the operator-valued covariance of $SECT(\cdot)$ corresponding to distribution $\mathbb{P}^{(j)}$ (see Eq. (3.3.6)) and $\Xi_{\nu^*}^{(j)}(s, t)$ the real-valued covariance of process $SECT(\cdot)(\nu^*)$ corresponding to $\mathbb{P}^{(j)}$, for $j \in \{1, 2\}$. Hereafter, we assume the following.

Assumption 3. $C^{(1)}(\nu_1, \nu_2) = C^{(2)}(\nu_1, \nu_2)$, for all $\nu_1, \nu_2 \in \mathbb{S}^{d-1}$.

Eq. (3.3.7) and Assumption 3 imply $\Xi_{\nu^*}^{(1)} = \Xi_{\nu^*}^{(2)} \stackrel{\text{def}}{=} \Xi_{\nu^*}$. The second half of Theorem 3.3.4 implies $\Xi_{\nu^*}(\cdot, \cdot) \in L^2([0, T]^2)$. Hence, we may define the integral operator

$$L^2([0, T]) \rightarrow L^2([0, T]), \quad f \mapsto \int_0^T f(s) \Xi_{\nu^*}(s, \cdot) ds. \quad (3.4.3)$$

Theorems VI.22 (e) and VI.23 of Reed (2012) indicate that the integral operator defined in Eq. (3.4.3) is compact. It is straightforward that this integral operator is also self-adjoint. Furthermore, the Hilbert-Schmidt theorem (Theorem VI.16 of Reed (2012)) implies that this operator has countably many eigenfunctions $\{\phi_l\}_{l=1}^{\infty}$ and eigenvalues $\{\lambda_l\}_{l=1}^{\infty}$ with $\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_l \geq \dots \geq 0$, and $\{\phi_l\}_{l=1}^{\infty}$ form an orthonormal basis of $L^2([0, T])$. Then, Theorems 3.3.3 and 3.3.4 imply the following expansions of the SECT in the distinguishing direction ν^* .

Theorem 3.4.1. (Karhunen–Loève expansion) For each $j \in \{1, 2\}$, suppose the random shapes $\{K_i^{(j)}\}_{i=1}^n \stackrel{i.i.d.}{\sim} \mathbb{P}^{(j)}$.

We have the following expansions for each sample $K_i^{(j)}$

$$SECT(K_i^{(j)})(\nu^*; t) = m_{\nu^*}^{(j)}(t) + \sum_{l=1}^{\infty} \sqrt{\lambda_l} \cdot Z_{l,i}^{(j)} \cdot \phi_l(t) \quad \text{for all } t \in [0, T], \quad (3.4.4)$$

$$\text{where } Z_{l,i}^{(j)} = \frac{1}{\sqrt{\lambda_l}} \int_0^T \left\{ SECT(K_i^{(j)})(\nu^*; t) - m_{\nu^*}^{(j)}(t) \right\} \phi_l(t) dt,$$

for all $l = 1, 2, \dots$, and the infinite series converges in the $L^2(\mathcal{S}_{R,d}^M, \mathbb{P}^{(j)})$ sense. For each fixed $j \in \{1, 2\}$ and $i \in \{1, \dots, n\}$, random variables $\{Z_{l,i}^{(j)}\}_{l=1}^{\infty}$ are defined on the probability space $(\mathcal{S}_{R,d}^M, \mathcal{F}, \mathbb{P}^{(j)})$, mutually uncorrelated, and have mean 0 and variance 1. Furthermore, since $\mathbb{P}^{(1)}$ and $\mathbb{P}^{(2)}$ are independent, the two collections $\{(Z_{l,1}^{(1)}, \dots, Z_{l,n}^{(1)})\}_{l=1}^{\infty}$ and $\{(Z_{l,1}^{(2)}, \dots, Z_{l,n}^{(2)})\}_{l=1}^{\infty}$ are independent.

We omit the proof of Theorem 3.4.1 as it follows directly from Corollary 5.5 of Alexanderian (2015). The uncorrelation of $\{Z_{l,i}^{(j)}\}_{l=1}^\infty$ across l in Theorem 3.4.1 plays an important role in our hypothesis testing framework.

3.4.2 The Theoretical Foundation for Hypothesis Testing

Suppose we have two collections of random samples $\{K_i^{(j)}\}_{i=1}^n \stackrel{i.i.d.}{\sim} \mathbb{P}^{(j)}$, for $j \in \{1, 2\}$. In this subsection, we provide the theoretical foundation of using $K_i^{(j)}$ to test the hypotheses in Eq. (3.4.1).

The Karhunen–Loève expansions in Eq. (3.4.4) provide the following

$$\begin{aligned} X_i(t) &\stackrel{\text{def}}{=} SECT(K_i^{(1)})(\nu^*; t) - SECT(K_i^{(2)})(\nu^*; t) \\ &= \left\{ m_{\nu^*}^{(1)}(t) - m_{\nu^*}^{(2)}(t) \right\} + \sum_{l=1}^{\infty} \sqrt{2\lambda_l} \cdot \left(\frac{Z_{l,i}^{(1)} - Z_{l,i}^{(2)}}{\sqrt{2}} \right) \cdot \phi_l(t). \end{aligned}$$

We further define the random variables $\{(\xi_{l,1}, \dots, \xi_{l,n})\}_{l=1}^\infty$ as follows

$$\begin{aligned} \xi_{l,i} &\stackrel{\text{def}}{=} \frac{1}{\sqrt{2\lambda_l}} \cdot \int_0^T X_i(t) \phi_l(t) dt \stackrel{(1)}{=} \theta_l + \left(\frac{Z_{l,i}^{(1)} - Z_{l,i}^{(2)}}{\sqrt{2}} \right), \\ \text{where } \theta_l &= \frac{1}{\sqrt{2\lambda_l}} \int_0^T \left\{ m_{\nu^*}^{(1)}(t) - m_{\nu^*}^{(2)}(t) \right\} \phi_l(t) dt \end{aligned} \tag{3.4.5}$$

and $\stackrel{(1)}{=}$ follows from that $\{\phi_l\}_{l=1}^\infty$ is an orthonormal basis of $L^2([0, T])$. Then, each $\xi_{l,i}$ has mean θ_l and variance 1. The following theorem transforms the null H_0 in Eq. (3.4.1) using the means $\{\theta_l\}_{l=1}^\infty$.

Theorem 3.4.2. *The null H_0 in Eq. (3.4.1) is equivalent to $\theta_l = 0$ for all $l = 1, 2, 3, \dots$.*

Proof. We have shown that the null H_0 is equivalent to $m_{\nu^*}^{(1)}(t) = m_{\nu^*}^{(2)}(t)$ for all $t \in [0, T]$, where ν^* is defined in Eq. (3.4.2). The null H_0 directly implies that $\theta_l = 0$ for all l . On the other hand, if $\theta_l = 0$ for all l , that $\{\phi_l\}_l$ is an orthonormal basis of $L^2([0, T])$ indicates that $m_{\nu^*}^{(1)} = m_{\nu^*}^{(2)}$ almost everywhere with respect to the Lebesgue measure dt . Theorem 3.3.3 (i) and the embedding $\mathcal{H} \subset \mathcal{B}$ in Eq. (3.1.2) imply that $m_{\nu^*}^{(1)}$ and $m_{\nu^*}^{(2)}$ are continuous functions. Then, $m_{\nu^*}^{(1)}(t) = m_{\nu^*}^{(2)}(t)$ for all $t \in [0, T]$, which is equivalent to the null H_0 in Eq. (3.4.1). $\square$

When eigenvalues λ_l in the denominators of Eq. (3.4.5) are close to zero for large l , the estimated θ_l corresponding to the small eigenvalues can be unstable. Specifically, even if $m_{\nu^*}^{(1)}(t) \approx m_{\nu^*}^{(2)}(t)$ for all t , an extremely small λ_l can move the corresponding θ_l far away from zero. Motivated by the cumulative variance $\int_0^T \mathbb{V}\{X_i(t)\} dt = 2(\sum_{l=1}^\infty \lambda_l)$ and the standard approach in principal component analysis, we focus on $\{\theta_l\}_{l=1}^L$ with

$$L \stackrel{\text{def}}{=} \min \left\{ l \in \mathbb{N} \mid \frac{\sum_{l'=1}^l \lambda_{l'}}{\sum_{l''=1}^\infty \lambda_{l''}} > 0.95 \right\}. \tag{3.4.6}$$

Hence, to test the null and alternative hypotheses in Eq. (3.4.1), we test the following

$$\begin{aligned}\widehat{H}_0 : \theta_1 = \theta_2 = \cdots = \theta_L = 0, \quad vs. \\ \widehat{H}_1 : \text{there exists } l' \in \{1, \dots, L\} \text{ such that } \theta_{l'} \neq 0.\end{aligned}\tag{3.4.7}$$

To test the hypotheses in Eq. (3.4.7), we need the following independence assumption.

Assumption 4. $\{(Z_{l,i}^{(1)} - Z_{l,i}^{(2)})/\sqrt{2}\}_{i=1}^n$ is independent of $\{(Z_{l',i}^{(1)} - Z_{l',i}^{(2)})/\sqrt{2}\}_{i=1}^n$ for any $l, l' \in \{1, \dots, L\}$ with $l \neq l'$.

Assumption 4 is satisfied under some conditions. The following theorem gives an example.

Theorem 3.4.3. For any partition of $[0, T]$, say $0 = t_0 < t_1 < \cdots < t_l = T$, the stochastic processes $\{\chi_t^{\nu^*} : t_0 \leq t \leq t_1\}, \dots, \{\chi_t^{\nu^*} : t_{l-1} \leq t \leq t_l\}$ are independent according to $\mathbb{P}^{(j)}$.

Proof. The continuity of the paths of $PECT(\cdot)(\nu^*) = \{\int_0^t \chi_\tau^{\nu^*} d\tau\}_{t \in [0, T]}$, together with $\int_0^0 \chi_\tau^{\nu^*} d\tau = 0$, implies that $PECT(\cdot)(\nu^*)$ is a Gaussian process according to $\mathbb{P}^{(j)}$ (see Theorem 14.4 of Kallenberg (2021)). Eq. (3.2.11) implies that $SECT(\cdot)(\nu^*)$ is a Gaussian process according to $\mathbb{P}^{(j)}$ as well. Then, for each fixed $i \in \{1, \dots, n\}$ and $j \in \{1, 2\}$, the $\{Z_{l,i}^{(j)}\}_{l=1}^\infty$ in Eq. (3.4.1) are i.i.d. standard normal random variables (see the discussion after Corollary 5.5 in Alexanderian (2015)). Therefore, Assumption 4 is satisfied. $\square$

Here, we provide simulations supporting Assumption 4. Suppose that $\mathbb{P}^{(1)}$ and $\mathbb{P}^{(2)}$ are equal to the distribution generating the random shapes defined in Eq. (3.3.9). We generate $\{K_i^{(j)}\}_{i=1}^n \stackrel{i.i.d.}{\sim} \mathbb{P}^{(j)}$ with $n = 100$ and 300 , for $j \in \{1, 2\}$. Using the numerical method proposed in Section 3.4.3, we approximately compute $\{\xi_{l,i}\}_{i=1}^n$, for $l \in \{1, 2\}$ (the L defined in Eq. (3.4.6) throughout all our simulations are 2 or 3). The mean functions $m_{\nu^*}^{(j)}$, for $j \in \{1, 2\}$, and the histograms of $\{\xi_{l,i}\}_{i=1}^n$ are presented in Figure 3.3. The overlapping curves of mean functions (blue solid and orange dashed curves in Figure 3.3) indicate that the null H_0 in Eq. (3.4.1) is true; hence, $\xi_{l,i} = (Z_{l,i}^{(1)} - Z_{l,i}^{(2)})/\sqrt{2}$ for $l \in \{1, 2\}$. More importantly, the normality of $\{\xi_{l,i}\}_{i=1}^n$ presented by the histograms in Figure 3.3 indicates that $\{\xi_{1,i}\}_{i=1}^n = \{(Z_{1,i}^{(1)} - Z_{1,i}^{(2)})/\sqrt{2}\}_{i=1}^n$ and $\{\xi_{2,i}\}_{i=1}^n = \{(Z_{2,i}^{(1)} - Z_{2,i}^{(2)})/\sqrt{2}\}_{i=1}^n$ are independent, instead of just uncorrelated (see Theorem 3.4.1). Furthermore, the simulation with $n = 300$ presents better normality and weaker correlation than the one with $n = 100$.

When L is very large, we may implement the following adaptive Neyman test statistic proposed in Fan (1996) to test the hypotheses in Eq. (3.4.7)

$$\begin{aligned}T_{L,i}^{(AN)} = & \sqrt{2 \log \log L} \cdot \max_{1 \leq l \leq L} \left\{ \frac{1}{\sqrt{2l}} \sum_{l'=1}^l (\xi_{l',i}^2 - 1) \right\} \\ & - \left\{ 2 \log \log L + \frac{1}{2} \log \log \log L - \frac{1}{2} \log(4\pi) \right\}.\end{aligned}\tag{3.4.8}$$

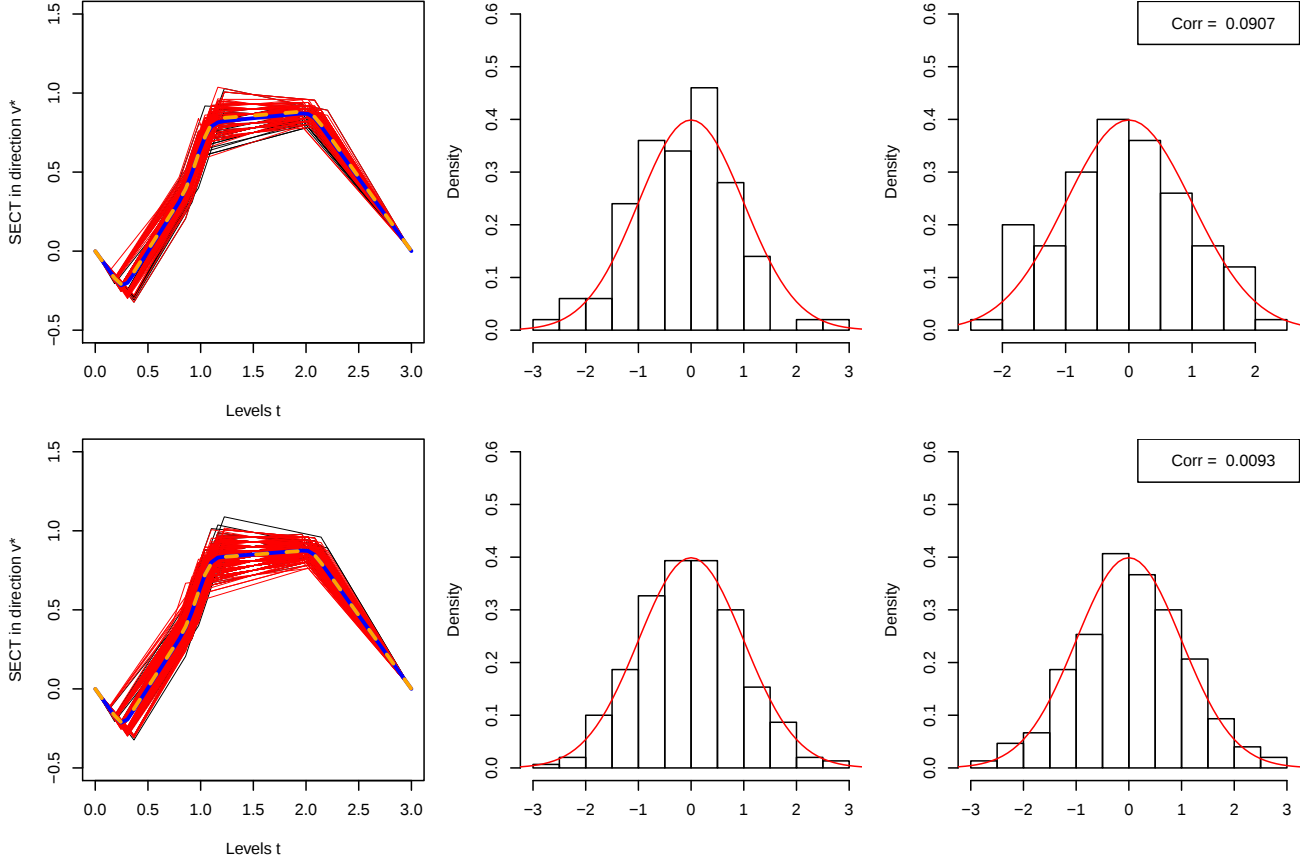


Figure 3.3: The first row presents the simulation results with $n = 100$ and the second row presents the simulation results with $n = 300$. In the left panel of each row, the back and red solid curves present the sample paths of $\{SECT(K_i^{(1)})(\nu^*)\}$ and $\{SECT(K_i^{(2)})(\nu^*)\}$, respectively, in the distinguishing direction ν^* (the red curves are thinner than the black ones and drawn on top of the black curves); the blue solid and red dashed curves present mean functions $m_{\nu^*}^{(1)}$ and $m_{\nu^*}^{(2)}$, respectively. The histograms of $\{\xi_{l,i}\}_{i=1}^n$, for $l \in \{1, 2\}$, in the middle and right panels are presented with the $N(0, 1)$ -density curve (red) overlaid on top of them. When $n = 100$, the sample correlation between $\{\xi_{1,i}\}_{i=1}^n$ and $\{\xi_{2,i}\}_{i=1}^n$ across $i \in \{1, \dots, n\}$ is 0.0907; when $n = 300$, this correlation is 0.00903.

Theorem 1 of [Darling and Erdős \(1956\)](#) implies that $\lim_{L \rightarrow \infty} \mathbb{P}(T_{L,i}^{(AN)} < t) = \exp\{-\exp(-t)\}$ under the null hypothesis and Assumption 4 (also see the discussions right after the Theorem 2.1 in [Fan \(1996\)](#)). Then $\hat{H}_0$ in Eq. (3.4.7) is rejected if $\{T_{L,i}^{(AN)}\}_{i=1}^n$ deviates from the distribution with the cumulative density function $\exp\{-\exp(-t)\}$. However, the number L defined in Eq. (3.4.6) is usually small. For example, if $\mathbb{P}^{(1)}$ and $\mathbb{P}^{(2)}$ are equal to the distribution generating random shapes defined in Eq. (3.3.9), then the L is equal to 2 or 3 across thousands of simulations. Hence, we implement a simple testing approach when L is not large, and we will focus on the simple testing approach for small L scenarios.

For each $l \in \{1, \dots, L\}$, the central limit theorem indicates that $\frac{1}{\sqrt{n}} \sum_{i=1}^n \xi_{l,i}$ asymptotically follows a $N(\theta_l, 1)$ distribution as $n \rightarrow \infty$. Under Assumption 4, $\{\frac{1}{\sqrt{n}} \sum_{i=1}^n \xi_{l,i}\}_{l=1}^L$ are independent across l . Even if Assumption 4 is violated, the uncorrelation property from Theorem 3.4.1 and the asymptotic normality of $\frac{1}{\sqrt{n}} \sum_{i=1}^n \xi_{l,i}$ provide the asymptotic independence of $\{\frac{1}{\sqrt{n}} \sum_{i=1}^n \xi_{l,i}\}_{l=1}^L$ across l . In either case, $\sum_{l=1}^L (\frac{1}{\sqrt{n}} \sum_{i=1}^n \xi_{l,i})^2$ is asymptotically χ_L^2 distributed under the null $\hat{H}_0$ in Eq. (3.4.7) as $n \rightarrow \infty$. Therefore, at the asymptotic confidence level $1 - \alpha$ for

$\alpha \in (0, 1)$, we reject $\hat{H}_0$ if

$$\sum_{l=1}^L \left(\frac{1}{\sqrt{n}} \sum_{i=1}^n \xi_{l,i} \right)^2 > \chi_{L,1-\alpha}^2 = \text{the } 1 - \alpha \text{ lower quantile of the } \chi_L^2 \text{ distribution.} \quad (3.4.9)$$

In Section 3.5, we will provide simulation studies for this testing procedure.

3.4.3 The Numerical foundation for Hypothesis Testing

In Section 3.4.2, we proposed an approach to testing the hypotheses in Eq. (3.4.1) based on $\{\xi_{l,i}\}_l$ defined in Eq. (3.4.5). In applications, neither the mean function $m_{\nu}^{(j)}(t)$ nor the covariance function $\Xi_{\nu}(t', t)$ is known. Hence, the corresponding Karhunen–Loève expansions in Eq. (3.4.4) are not available, and the proposed hypothesis testing approach is not directly applicable. In this subsection, motivated by Chapter 4.3.2 of Rasmussen and Williams (2006), we propose a method for estimating the $\{\xi_{l,i}\}_l$ in Eq. (3.4.5), which enables us to conduct statistical inference.

For random shapes $\{K_i^{(j)}\}_{i=1}^n \stackrel{i.i.d.}{\sim} \mathbb{P}^{(j)}$ coming from the two distributions $j \in \{1, 2\}$, we compute each of their SECT in finitely many directions and at finitely many levels as discussed in Section 3.3.2 and get $\{SECT(K_i^{(j)})(\nu_p; t_q) \mid p = 1, \dots, \Gamma \text{ and } q = 1, \dots, \Delta\}_{i=1}^n$ for $j \in \{1, 2\}$, where $t_q = \frac{T}{\Delta}q$. We want to emphasize that the SECT of all shapes $K_i^{(j)}$ in the two collections are computed in the same collection of directions $\{\nu_p\}_{p=1}^{\Gamma}$ and at the same collection of levels $\{t_q\}_{q=1}^{\Delta}$. We estimate the mean $m_{\nu_p}^{(j)}(t_q)$ at level t_q by $\hat{m}_{\nu_p}^{(j)}(t_q)$ = the sample mean of $\{SECT(K_i^{(j)})(\nu_p; t_q)\}_{i=1}^n$ across $i \in \{1, \dots, n\}$. Then, we estimate the distinguishing direction ν^* by

$$\hat{\nu}^* \stackrel{\text{def}}{=} \arg \max_{\nu_p} \left[\max_{t_q} \left\{ \left| \hat{m}_{\nu_p}^{(1)}(t_q) - \hat{m}_{\nu_p}^{(2)}(t_q) \right| \right\} \right]. \quad (3.4.10)$$

Based on Assumption 3, we estimate the covariance matrix $(\Xi_{\nu^*}(t_{q'}, t_q))_{q, q'=1, \dots, \Delta}$ by the sample covariance matrix $\mathbf{C} = (\hat{\Xi}_{\nu^*}(t_{q'}, t_q))_{q, q'=1, \dots, \Delta}$ derived from the following centered sample vectors across $i \in \{1, \dots, n\}$ and $j \in \{1, 2\}$

$$\left\{ \left(SECT(K_i^{(j)})(\hat{\nu}^*; t_1) - \hat{m}_{\hat{\nu}^*}^{(j)}(t_1), \dots, SECT(K_i^{(j)})(\hat{\nu}^*; t_{\Delta}) - \hat{m}_{\hat{\nu}^*}^{(j)}(t_{\Delta}) \right)^{\top} \mid j = 1, 2 \right\}_{i=1}^n. \quad (3.4.11)$$

Since the eigenfunctions $\{\phi_l\}_{l=1}^{\infty}$ and eigenvalues $\{\lambda_l\}_{l=1}^{\infty}$ satisfy $\lambda_l \phi_l = \int_0^T \phi_l(s) \Xi_{\nu^*}(s, \cdot) ds$, we have the following approximation

$$\lambda_l \phi_l(t_q) = \int_0^T \phi_l(s) \Xi_{\nu^*}(s, t_q) ds \approx \frac{T}{\Delta} \sum_{q'=1}^{\Delta} \phi_l(t_{q'}) \Xi_{\nu^*}(t_{q'}, t_q) \approx \frac{T}{\Delta} \sum_{q'=1}^{\Delta} \phi_l(t_{q'}) \hat{\Xi}_{\nu^*}(t_{q'}, t_q),$$

which is represented in the following matrix form

$$\lambda_l \begin{pmatrix} \phi_l(t_1) \\ \vdots \\ \phi_l(t_\Delta) \end{pmatrix} \approx \frac{T}{\Delta} \begin{pmatrix} \widehat{\Xi}_{\nu^*}(t_1, t_1) & \dots & \widehat{\Xi}_{\nu^*}(t_\Delta, t_1) \\ \vdots & \ddots & \vdots \\ \widehat{\Xi}_{\nu^*}(t_\Delta, t_2) & \dots & \widehat{\Xi}_{\nu^*}(t_\Delta, t_\Delta) \end{pmatrix} \begin{pmatrix} \phi_l(t_1) \\ \vdots \\ \phi_l(t_\Delta) \end{pmatrix}. \quad (3.4.12)$$

We denote the eigenvectors and eigenvalues of $\mathbf{C}$ as $\{\mathbf{v}_l = (v_{l,1}, \dots, v_{l,\Delta})^\top\}_{l=1}^\Delta$ and $\{\Lambda_l\}_{l=1}^\Delta$, respectively. The following equation motivates the estimate $\phi_l(t_q) \approx \widehat{\phi}_l(t_q) \stackrel{\text{def}}{=} \sqrt{\frac{\Delta}{T}} \cdot v_{l,q}$, for all $l \in \{1, \dots, \Delta\}$

$$\sum_{q=1}^\Delta v_{l,q}^2 = \|\mathbf{v}_l\|^2 = 1 = \int_0^T |\phi_l(t)|^2 dt \approx \frac{T}{\Delta} \sum_{q=1}^\Delta (\phi_l(t_q))^2 = \sum_{q=1}^\Delta \left(\sqrt{\frac{T}{\Delta}} \cdot \phi_l(t_q) \right)^2.$$

The following equation motivates the estimate $\lambda_l \approx \widehat{\lambda}_l \stackrel{\text{def}}{=} \frac{T}{\Delta} \Lambda_l$, for all $l \in \{1, \dots, \Delta\}$

$$\begin{aligned} \lambda_l \left(\widehat{\phi}_l(t_1), \dots, \widehat{\phi}_l(t_\Delta) \right)^\top &\approx \frac{T}{\Delta} \mathbf{C} \left(\widehat{\phi}_l(t_1), \dots, \widehat{\phi}_l(t_\Delta) \right)^\top \\ &= \sqrt{\frac{T}{\Delta}} \mathbf{C} \mathbf{v}_l = \sqrt{\frac{T}{\Delta}} \Lambda_l \mathbf{v}_l = \left(\frac{T}{\Delta} \Lambda_l \right) \left(\widehat{\phi}_l(t_1), \dots, \widehat{\phi}_l(t_\Delta) \right)^\top. \end{aligned}$$

Additionally, we estimate the L defined in Eq. (3.4.6) by the following

$$L \approx \widehat{L} \stackrel{\text{def}}{=} \min \left\{ l = 1, \dots, \Delta \left| \frac{\sum_{l'=1}^l |\widehat{\lambda}_{l'}|}{\sum_{l''=1}^\Delta |\widehat{\lambda}_{l''}|} > 0.95 \right. \right\}; \quad (3.4.13)$$

we take the absolute values of the estimated eigenvalues in computing $\widehat{L}$ because the estimated eigenvalues may be numerically negative in applications.

We estimate the random variables $\xi_{l,i}$ defined in Eq. (3.4.5) by the following

$$\xi_{l,i} \approx \widehat{\xi}_{l,i} = \frac{1}{\sqrt{2\widehat{\lambda}_l}} \cdot \frac{T}{\Delta} \sum_{q=1}^\Delta \left\{ SECT(K_i^{(1)})(\widehat{\nu}^*; t_q) - SECT(K_i^{(2)})(\widehat{\nu}^*; t_q) \right\} \widehat{\phi}_l(t_q), \quad (3.4.14)$$

for $l = 1, \dots, \widehat{L}$ and $i = 1, \dots, n$. Then, when $\widehat{L}$ is large (e.g., several hundreds; see Section 3 of Fan (1996)), we implement the adaptive Neyman test by replacing the $\xi_{l,i}$ and L in Eq. (3.4.8) with $\widehat{\xi}_{l,i}$ and $\widehat{L}$, respectively. When $\widehat{L}$ is moderate, which is true in all our simulations, we can implement the χ^2 -test in Eq. (3.4.9) as follows

$$\sum_{l=1}^{\widehat{L}} \left(\frac{1}{\sqrt{n}} \sum_{i=1}^n \widehat{\xi}_{l,i} \right)^2 > \chi_{\widehat{L}, 1-\alpha}^2 = \text{the } 1-\alpha \text{ lower quantile of the } \chi_{\widehat{L}}^2 \text{ distribution.} \quad (3.4.15)$$

A complete outline of this χ^2 -hypothesis testing procedure is given in Algorithm 4.

In addition to the χ^2 -test detailed in Algorithm 4, we also propose a permutation test as an alternative approach

Algorithm 4 : (χ^2 -based) Testing the hypotheses in Eq. (3.4.1)

Input: (i) SECT of two collection of shapes $\{SECT(K_i^{(j)})(\nu_p; t_q) : p = 1, \dots, \Gamma \text{ and } q = 1, \dots, \Delta\}_{i=1}^n$ for $j \in \{1, 2\}$; (ii) desired confidence level $1 - \alpha$ with $\alpha \in (0, 1)$.

Output: **Accept** or **Reject** the null hypothesis H_0 in Eq. (3.4.1).

- 1: For $j \in \{1, 2\}$, compute $\hat{m}_{\nu_p}^{(j)}(t_q) \stackrel{\text{def}}{=} \text{sample mean of } \{SECT(K_i^{(j)})(\nu_p; t_q)\}_{i=1}^n$ across $i \in \{1, \dots, n\}$.
 - 2: Compute the estimated distinguishing direction $\hat{\nu}^*$ using Eq. (3.4.10).
 - 3: Compute $\mathbf{C} = (\hat{\Xi}_{\nu^*}(t_{q'}, t_q))_{q, q'=1, \dots, \Delta} \stackrel{\text{def}}{=}$ the sample covariance matrix derived from the centered sample vectors in Eq. (3.4.11) across $i \in \{1, \dots, n\}$ and $j \in \{1, 2\}$.
 - 4: Compute the eigenvectors $\{\mathbf{v}_l\}_{l=1}^{\Delta}$ and eigenvalues $\{\Lambda_l\}_{l=1}^{\Delta}$ of the sample covariance matrix $\mathbf{C}$.
 - 5: Compute $\hat{\phi}_l(t_q) \stackrel{\text{def}}{=} \sqrt{\frac{\Delta}{T}} v_{l,q}$ and $\hat{\lambda}_l \stackrel{\text{def}}{=} \frac{T}{\Delta} \Lambda_l$ for all $l = 1, \dots, \Delta$.
 - 6: Compute $\hat{L}$ using Eq. (3.4.13).
 - 7: Compute $\{\xi_{l,i} : l = 1, \dots, \Delta\}_{i=1}^n$ using Eq. (3.4.14), test null H_0 using Eq. (3.4.15), and report the output.
-

for assessing the statistical hypotheses in Eq. (3.4.1). The main idea behind the permutation test is that, under the null hypothesis, shuffling (or permuting) the group labels of shapes should not heavily change the test statistic of interest. To perform the permutation test, we first apply Algorithm 4 to our original data and then repeatedly re-apply Algorithm 4 to shapes with shuffled labels.[§] We then compare how the test statistics derived from the original differ from those computed on the shuffled data. The details of this permutation-based approach are provided in Algorithm 5. Simulation studies in Section 3.5 show that Algorithm 5 can eliminate the moderate type I error inflation of Algorithm 4; however, the power under the alternative for Algorithm 5 is moderately weaker than that of Algorithm 4.

Algorithm 5 : (Permutation-based) Testing the hypotheses in Eq. (3.4.1)

Input: (i) SECT of two collection of shapes $\{SECT(K_i^{(j)})(\nu_p; t_q) : p = 1, \dots, \Gamma \text{ and } q = 1, \dots, \Delta\}_{i=1}^n$ for $j \in \{1, 2\}$; (ii) desired confidence level $1 - \alpha$ with $\alpha \in (0, 1)$; (iii) the number of permutations Π .

Output: **Accept** or **Reject** the null hypothesis H_0 in Eq. (3.4.1).

- 1: Apply Algorithm 4 to the original input SECT data, compute $\hat{L}_0$ using Eq. (3.4.13) (see footnote §), and compute the χ^2 -test statistic denoted as $\mathfrak{S}_0$ using Eq. (3.4.15).
 - 2: **for all** $k = 1, \dots, \Pi$, **do**
 - 3: Randomly permute the group labels $j \in \{1, 2\}$ of the input SECT data.
 - 4: Apply Algorithm 4 to the permuted SECT data while setting $\hat{L}$ to be the $\hat{L}_0$, instead of using Eq. (3.4.13), and compute a χ^2 -test statistic $\mathfrak{S}_k$ using Eq. (3.4.15).
 - 5: **end for**
 - 6: Compute $k^* \stackrel{\text{def}}{=} \lceil (1 - \alpha) \cdot \Pi \rceil \stackrel{\text{def}}{=} \text{the largest integer smaller than } (1 - \alpha) \cdot \Pi$.
 - 7: **Reject** the null hypothesis H_0 if $\mathfrak{S}_0 > \mathfrak{S}_{k^*}$ and report the output.
-

3.5 Experiments Using Simulations

In this section, we present proof-of-concept simulation studies showing the performance of the hypothesis testing framework outlined in Algorithms 4 and 5. We focus on a family of distributions $\{\mathbb{P}^{(\varepsilon)}\}_{0 \leq \varepsilon \leq 0.1}$ on $\mathcal{S}_{R,d}^M$. For each

§: When we apply Algorithm 4 to the original SECT, the result of Eq. (3.4.13) is denoted as $\hat{L}_0$. When we apply Algorithm 4 to the shuffled SECT, the $\hat{L}$ resulting from Eq. (3.4.13) may differ from $\hat{L}_0$. To make the comparison between $\mathfrak{S}_0$ and $\mathfrak{S}_{k^*}$ fair (see the last step of Algorithm 5), we set $\hat{L}$ to be $\hat{L}_0$.

ε , a collection of shapes $\{K_i^{(\varepsilon)}\}_{i=1}^n$ are i.i.d. generated from distribution $\mathbb{P}^{(\varepsilon)}$ via the following

$$\begin{aligned} K_i^{(\varepsilon)} &= \left\{ x \in \mathbb{R}^2 \left| \inf_{y \in S_i^{(\varepsilon)}} \|x - y\| \leq \frac{1}{5} \right. \right\}, \quad \text{where} \\ S_i^{(\varepsilon)} &= \left\{ \left(\frac{2}{5} + a_{1,i} \cdot \cos t, b_{1,i} \cdot \sin t \right) \left| \frac{1-\varepsilon}{5}\pi \leq t \leq \frac{9+\varepsilon}{5}\pi \right. \right\} \\ &\quad \cup \left\{ \left(-\frac{2}{5} + a_{2,i} \cdot \cos t, b_{2,i} \cdot \sin t \right) \left| \frac{6\pi}{5} \leq t \leq \frac{14\pi}{5} \right. \right\}, \end{aligned} \quad (3.5.1)$$

where $\{a_{1,i}, a_{2,i}, b_{1,i}, b_{2,i}\}_{i=1}^n \stackrel{i.i.d.}{\sim} N(1, 0.05^2)$, and the index ε denotes the dissimilarity between the distributions $\mathbb{P}^{(\varepsilon)}$ and $\mathbb{P}^{(0)}$. For each $\varepsilon \in [0, 0.1]$, we test the following hypotheses using the scheme described in Algorithms 4 and 5

$$\begin{aligned} H_0 : m_\nu^{(0)}(t) &= m_\nu^{(\varepsilon)}(t) \text{ for all } (\nu, t) \in \mathbb{S}^{d-1} \times [0, T], \\ \text{vs. } H_1 : m_\nu^{(0)}(t) &\neq m_\nu^{(\varepsilon)}(t) \text{ for some } (\nu, t), \end{aligned} \quad (3.5.2)$$

where the mean $m_\nu^{(\varepsilon)}(t) \stackrel{\text{def}}{=} \int_{\mathcal{S}_{R,d}^M} SECT(K)(\nu; t) \mathbb{P}^{(\varepsilon)}(dK)$, and the null hypothesis H_0 is true when $\varepsilon = 0$.

We set $T = 3$, directions $\nu_p = (\cos \frac{p-1}{4}\pi, \sin \frac{p-1}{4}\pi)^\top$ for $p \in \{1, 2, 3, 4\}$, levels $t_q = \frac{T}{50}q = 0.06q$ for $q \in \{1, \dots, 50\}$, the confidence level 95% (i.e., $\alpha = 0.05$ in Eq. (3.4.15)), and the number of permutations $\Pi = 1000$. For each $\varepsilon \in \{0, 0.0125, 0.025, 0.0375, 0.05, 0.075, 0.1\}$, we independently generate two collection of shapes: $\{K_i^{(0)}\}_{i=1}^n \stackrel{i.i.d.}{\sim} \mathbb{P}^{(0)}$ and $\{K_i^{(\varepsilon)}\}_{i=1}^n \stackrel{i.i.d.}{\sim} \mathbb{P}^{(\varepsilon)}$ through Eq (3.5.1) with $n = 100$, and we compute the SECT of each generated shape in directions $\{\nu_p\}_{p=1}^4$ and at levels $\{t_q\}_{q=1}^{50}$. Then, we implement Algorithms 4 and 5 to these computed SECT variables and get the corresponding **Accept/Reject** output. We repeat this procedure for 100 times, and the rejection rates across all replicates for each ε are presented in Figure 3.4 and Table 3.1.

Table 3.1: Indices ε and the rejection rates of Algorithms 4 and 5.

Indices ε	0.000	0.0125	0.0250	0.0375	0.0500	0.0750	0.100
Rejection rates of Algorithm 4	0.15	0.16	0.32	0.67	0.92	0.98	1.00
Rejection rates of Algorithm 5	0.03	0.14	0.22	0.47	0.74	1.00	1.00

The presented simulation studies show that our Algorithms 4 and 5 are extremely powerful in detecting the difference between $\mathbb{P}^{(\varepsilon)}$ and $\mathbb{P}^{(0)}$ in terms of the discrepancy between the corresponding mean functions. The estimated mean functions $\hat{m}_{\hat{\nu}^*}^{(0)}$ and $\hat{m}_{\hat{\nu}^*}^{(\varepsilon)}$ in direction $\hat{\nu}^*$ are presented by the blue solid and red dashed curves, respectively, in Figure 3.4. The discrepancy between the two estimated mean functions is presented therein. As ε increases, $\mathbb{P}^{(\varepsilon)}$ deviates from $\mathbb{P}^{(0)}$, and the power of Algorithms 4 and 5 — rejection rates under the alternative hypothesis — in detecting the deviation increases. When $\varepsilon > 0.075$, the power of Algorithms 4 and 5 exceeds 0.95. For all values of ε , it is hard to tell the deviation of $\mathbb{P}^{(\varepsilon)}$ from $\mathbb{P}^{(0)}$ by eye. For example, by just visualizing the

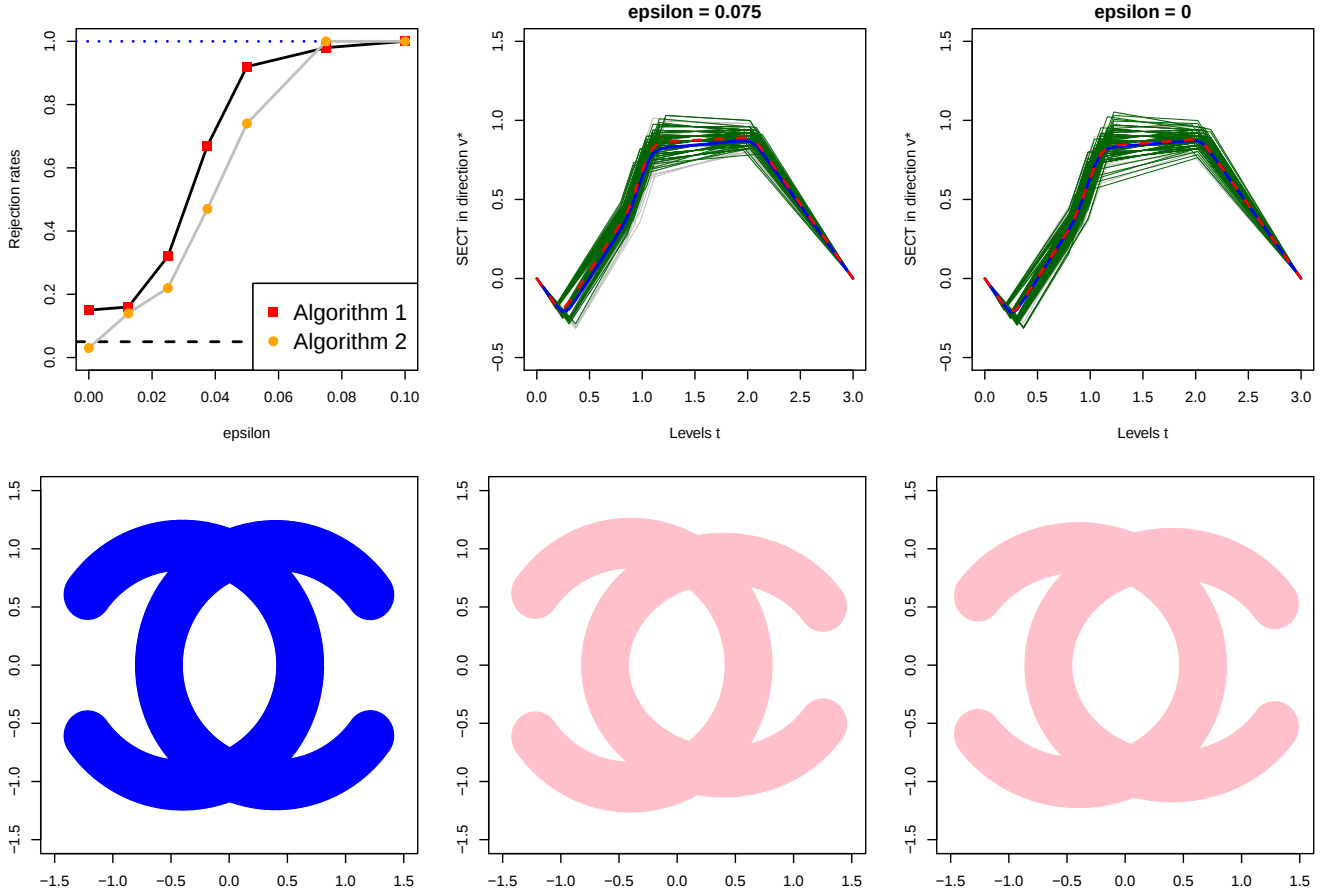


Figure 3.4: The “rejection rates vs. ϵ ” plot in the first panel presents the relationship between the indices ϵ and the rejection rates computed via Algorithms 4 and 5 throughout all simulations. This relationship is presented in Table 3.1 as well. The red squares and orange dots present the rejection rates corresponding to values $\epsilon \in \{0, 0.0125, 0.025, 0.0375, 0.05, 0.075, 0.1\}$. The black dashed and blue dotted straight lines present 0.05 and 1, respectively. In the last two panels of the first row, the gray and green solid curves present sample path collections $\{SECT(K_i^{(0)})(\hat{v}^*)\}_{i=1}^{100}$ and $\{SECT(K_i^{(\epsilon)})(\hat{v}^*)\}_{i=1}^{100}$, respectively; the blue solid and red dashed curves present $\hat{m}_{\hat{v}^*}^{(0)}$ and $\hat{m}_{\hat{v}^*}^{(\epsilon)}$, respectively, for $\epsilon \in \{0, 0.075\}$. The blue shape in the second row is generated from $\mathbb{P}^{(0)}$, and the two pink shapes are generated from $\mathbb{P}^{(0.075)}$.

shapes in Figure 3.4, it is difficult to distinguish between the shape collections generated by $\mathbb{P}^{(0)}$ (blue) and $\mathbb{P}^{(0.075)}$ (pink), while our hypothesis testing framework detected the difference between the two shape collections in more than 95% of the simulations.

Despite the strong ability of Algorithm 4 to detect the true discrepancy between mean functions, its type I error rate — rejection rate under the null model — is moderately inflated. Specifically, the type I error rate of Algorithm 4 is 0.15, while the expected error rate is 0.05 (see Table 3.1). In contrast, Algorithm 5 does not suffer from the type I error inflation under the null, while its power under the alternative is moderately weaker than that of Algorithm 4.

3.6 Conclusions and Discussions

In this paper, we provided the mathematical foundations for the randomness of shapes and the distributions of the smooth Euler characteristic transform (SECT). The $(\mathcal{S}_{R,d}^M, \mathcal{B}(\rho), \mathbb{P})$ was constructed as an underlying probability space for modeling the randomness of shapes, and the SECT was modeled as a $C(\mathbb{S}^{d-1}; \mathcal{H})$ -valued random variable defined on this probability space. We showed several properties of the SECT ensuring its Karhunen–Loève expansion, which led to normal distribution-based statistics $\sum_{l=1}^L (\frac{1}{\sqrt{n}} \sum_{i=1}^n \xi_{l,i})^2$ and $\{T_{L,i}^{(AN)}\}_{i=1}^n$ for testing the hypotheses on random shapes. Simulation studies were provided to support our mathematical derivations and show the performance of the proposed hypothesis testing framework. Our approach was shown to be extremely powerful in detecting the difference between the mean functions corresponding to two distributions of shapes. We list several potential future research areas that we believe are related to our work.

Definition of Mean Shapes and Fréchet/Wasserstein Regression The existence and uniqueness of the mean shapes $K_{\oplus}$ as defined in Eq. (3.3.10) are still unknown. If mean shapes $K_{\oplus}$ do exist, the relationship between $SECT(K_{\oplus})$ and $\mathbb{E}\{SECT(\cdot)\}$ is of particular interest. The Fréchet mean in Eq. (3.3.10) may be extended to the conditional Fréchet mean and implemented in the Fréchet regression — predicting shapes K using multiple scalar predictors (Petersen and Müller, 2019). For example, predicting molecular shapes and structures from scalar-valued indicators or sequences has become of high interest in systems biology (Jumper et al., 2021; Yang et al., 2020). As an example of the other way around, predicting clinical outcomes from the tumors is also of interest (Moon et al., 2020; Somasundaram et al., 2021; Vipond et al., 2021). Crawford et al. (2020) conducted this prediction using Gaussian process regression. Another possible way of making these types of predictions is via Wasserstein regression (Chen et al., 2021). The Wasserstein regression framework also allows both predictors and responses to be random shapes. For example, both the treatment of interest and the shapes of tumors pre-treatment can play the role of predictors, while tumor shapes after treatment can be used as the responses of a Wasserstein regression model. The combination of Wasserstein regression and the random shape framework proposed herein is left for future research.

Investigating the Gaussianity of the SECT Both the proof of Theorem 3.4.3 and the simulation results presented in Figure 3.3 indicate that modeling the real-valued stochastic process $SECT(\cdot)(\nu)$, for each $\nu \in \mathbb{S}^{d-1}$, as a Gaussian process is suitable for some underlying distributions $\mathbb{P}$ of shapes K . If $SECT(\cdot)(\nu)$ are Gaussian processes, we may further model the distributions of $SECT(\cdot)(\nu)$ using parametric covariance functions. Finding the proper assumptions making each $SECT(\cdot)(\nu)$ a Gaussian process is of both theoretical and applied modeling interest.

Generative Models for Complex Shapes The foundations for the randomness of shapes and the distributions of the SECT allow for the generative modeling of complex shapes. Suppose we are interested in a collection of shapes $\{K_i\}_{i=1}^n \subset \mathcal{S}_{R,d}^M$ (e.g., heel bones or tumors), and the underlying distribution of $\{SECT(K_i)\}_{i=1}^n$ is $\mathbf{P}$, where $\mathbf{P}$ is a probability measure on $C(\mathbb{S}^{d-1}; \mathcal{H})$. Suppose $\mathbf{P}$ is known or accurately estimated. We can then simulate $\{SECT_{i'}\}_{i'=1}^{n'} \stackrel{i.i.d.}{\sim} \mathbf{P}$. Since the shape-to-SECT map $K \mapsto SECT(K)$ is invertible (see Theorem 3.2.4), for each $i' \in \{1, \dots, n'\}$, there exists a unique shape $K_{i'}$ whose SECT is $SECT_{i'}$. Then the simulated $\{K_{i'}\}_{i'=1}^{n'}$ and the shapes of interest share the same distribution $\mathbf{P}$. The estimation of the underlying distribution $\mathbf{P}$ from $\{K_i\}_{i=1}^n$ is left for future research. The shape reconstruction step (i.e., constructing an (approximate) inverse map $SECT_i \mapsto K_i$) is outside the scope of the current paper; however, when shapes K are 3-dimensional “meshes,” a shape reconstruction approach can be applied (see Section 2.4 of Wang et al. (2021)).

Lifted Euler Characteristic Transform Kirveslahti and Mukherjee (2021) generalized the Euler characteristic transform to field type data using a lifting argument and proposed the lifted Euler characteristic transform (LECT). The randomness and distributions of the LECT and corresponding hypothesis testing framework extensions are left for future research.

Software Availability

Source code for implementing the simulation studies in Section 3.5 is publicly available online at

<https://github.com/KMengBrown/Randomness-Inference-via-SECT.git>.

Acknowledgments

We want to thank Dr. Matthew T. Harrison in the Division of Applied Mathematics at Brown University for useful comments and suggestions on previous versions of the manuscript.

Competing Interests

The authors declare no competing interests.

Appendix

3.A Overview of Persistent Diagrams

This subsection gives an overview of PDs in the literature. The overview is provided for the following purposes:

- we provide the details for the definition of $\mathcal{S}_{R,d}^M$, particularly the condition in Eq. (3.2.4);
- the PD framework is the necessary tool for proving several theorems in this paper (see Appendix 3.C).

Most of the materials in the overview come from or are modified from Mileyko et al. (2011) and Turner (2013).

Let $\mathbb{K}$ be a compact topological space and φ be a real-valued continuous function defined on $\mathbb{K}$. Because of the compactness of $\mathbb{K}$ and continuity of φ , we assume $\varphi(\mathbb{K}) \subset [0, T]$ without loss of generality. For each $t \in [0, T]$, denote

$$\mathbb{K}_t^\varphi \stackrel{\text{def}}{=} \{x \in \mathbb{K} \mid \varphi(x) \leq t\}.$$

Then $\mathbb{K}_{t_1}^\varphi \subset \mathbb{K}_{t_2}^\varphi$ for all $0 \leq t_1 \leq t_2 \leq T$, and $i_{t_1 \rightarrow t_2}$ denotes the corresponding inclusion map. Definition 3.2.1 is an analogue to the following concepts:

- i) t^* is a homotopy critical point (HCP) of $\mathbb{K}$ with respect to φ if $\mathbb{K}_{t^*}^\varphi$ is not homotopy equivalent to $\mathbb{K}_{t^*-\delta}^\varphi$ for any $\delta > 0$;
- ii) $\mathbb{K}$ is called tame with respect to φ if $\mathbb{K}$ has finitely many HCP with respect to φ .

If we take $\mathbb{K} = K \in \mathcal{S}_{R,d}^M$ and

$$\varphi(x) = x \cdot \nu + R \stackrel{\text{def}}{=} \phi_\nu(x), \quad x \in K, \quad \nu \in \mathbb{S}^{d-1}, \quad (3.A.1)$$

we have the scenario discussed in Section 3.2.1. The definition of $\mathcal{S}_{R,d}^M$ provides the tameness of K with respect to ϕ_ν and $\phi_\nu(K) \subset [0, T]$, for any fixed $\nu \in \mathbb{S}^{d-1}$.

The inclusion maps $i_{t_1 \rightarrow t_2} : \mathbb{K}_{t_1}^\varphi \rightarrow \mathbb{K}_{t_2}^\varphi$ induces the group homomorphisms

$$i_{t_1 \rightarrow t_2}^\# : H_k(\mathbb{K}_{t_1}^\varphi) \rightarrow H_k(\mathbb{K}_{t_2}^\varphi), \quad \text{for all } k \in \mathbb{Z},$$

where $H_k(\cdot) = H_k(\cdot; \mathbb{Z}_2)$ denotes the k -th homology group with respect to field $\mathbb{Z}_2$, and $\mathbb{Z}_2$ is omitted for simplicity.

Because of the tameness of $\mathbb{K}$ with respect to φ , for any $t_1 \leq t_2$, we have that the image

$$\text{im} \left(i_{(t_1-\delta) \rightarrow t_2}^\# \right) = \text{im} \left(i_{(t_1-\delta) \rightarrow t_1}^\# \circ i_{t_1 \rightarrow t_2}^\# \right)$$

does not depend on $\delta > 0$ when δ is sufficiently small, and then this constant image is denoted as $im(i_{(t_1-)\rightarrow t_2}^\#)$. For any t , the k -th birth group at t is defined as the quotient group

$$B_k^t \stackrel{\text{def}}{=} H_k(\mathbb{K}_t^\varphi)/im(i_{(t-)\rightarrow t}^\#),$$

and $\pi_{B_k^t} : H_k(\mathbb{K}_t^\varphi) \rightarrow B_k^t$ denotes the corresponding quotient map. For any $\alpha \in H_k(\mathbb{K}_t^\varphi)$, we say α is born at t if $\pi_{B_k^t}(\alpha) \neq 0$ in B_k^t . The tameness implies that B_k^t is a nontrivial group only for finitely many t . For any $t_1 < t_2$, we denote the quotient group

$$E_k^{t_1, t_2} \stackrel{\text{def}}{=} H_k(\mathbb{K}_{t_2}^\varphi)/im(i_{(t_1-)\rightarrow t_2}^\#)$$

and the corresponding quotient map $\pi_{E_k^{t_1, t_2}} : H_k(\mathbb{K}_{t_2}^\varphi) \rightarrow E_k^{t_1, t_2}$. Furthermore, we define the following map

$$g_k^{t_1, t_2} : B_k^{t_1} \rightarrow E_k^{t_1, t_2}, \quad \pi_{B_k^{t_1}}(\alpha) \mapsto \pi_{E_k^{t_1, t_2}}(i_{t_1 \rightarrow t_2}^\#(\alpha)),$$

for all $\alpha \in H_k(\mathbb{K}_{t_1}^\varphi)$. Then we define the death group

$$D_k^{t_1, t_2} \stackrel{\text{def}}{=} \ker(g_k^{t_1, t_2}).$$

We say a homology class $\alpha \in H_k(\mathbb{K}_{t_1}^\varphi)$ is born at t_1 and dies at t_2 if (i) $\pi_{B_k^{t_1}}(\alpha) \neq 0$, (ii) $\pi_{B_k^{t_1}}(\alpha) \in D_k^{t_1, t_2}$, and (iii) $\pi_{B_k^{t_1}}(\alpha) \notin D_k^{t_1, t_2-\delta}$ for any $\delta \in (0, t_2 - t_1)$. If α does not die, we artificially say that it dies at T as $K_T^\varphi = \mathbb{K}$. Then we denote $\text{birth}(\alpha) = t_1$ and $\text{death}(\alpha) = t_2$, and the persistence of α is defined as

$$\text{pers}(\alpha) \stackrel{\text{def}}{=} \text{death}(\alpha) - \text{birth}(\alpha).$$

With the notions of $\text{death}(\alpha)$ and $\text{birth}(\alpha)$, the k -th PD of $\mathbb{K}$ with respect to φ is defined as the following multiset of 2-dimensional points (see Definition 2 of [Mileyko et al. \(2011\)](#)).

$$\text{Dgm}_k(\mathbb{K}; \varphi) \stackrel{\text{def}}{=} \left\{ (\text{birth}(\alpha), \text{death}(\alpha)) \mid \alpha \in H_k(\mathbb{K}_t) \text{ for some } t \in [0, T] \text{ with } \text{pers}(\alpha) > 0 \right\} \bigcup \mathfrak{D}, \quad (3.A.2)$$

where $(\text{birth}(\alpha_1), \text{death}(\alpha_1))$ and $(\text{birth}(\alpha_2), \text{death}(\alpha_2))$ for $\alpha_1 \neq \alpha_2$ are counted as two points even if α_1 and α_2 are born and die at the same times, respectively; that is, the multiplicity of the point $(\text{birth}(\alpha_1), \text{death}(\alpha_1)) = (\text{birth}(\alpha_2), \text{death}(\alpha_2))$ is at least 2; $\mathfrak{D}$ denotes the diagonal $\{(t, t) \mid t \in \mathbb{R}\}$ with the multiplicity of each point on this diagonal is the cardinality of $\mathbb{Z}$. Since $\text{birth}(\alpha)$ is no later than $\text{death}(\alpha)$, the PD $\text{Dgm}_k(\mathbb{K}; \varphi)$ is contained in the triangular region $\{(s, t) \in \mathbb{R}^2 : 0 \leq s \leq t \leq T\}$.

Condition (3.2.4) in the definition of $\mathcal{S}_{R,d}^M$ The function ϕ_ν defined in Eq. (3.A.1), the corresponding PDs defined by Eq. (3.A.2), and the definition of $\text{pers}(\cdot)$ provide the details of condition (3.2.4). The notation $\#\{\cdot\}$ counts the multiplicity of the corresponding multiset.

Generally, a persistence diagram is a countable multiset of points in triangular region $\{(s, t) \in \mathbb{R}^2 : 0 \leq s, t \leq T \text{ and } s \leq t\}$ along with $\mathfrak{D}$ (see Definition 2 of Mileyko et al. (2011)). The collection of all persistence diagrams is denoted as $\mathcal{D}$. Obviously, all the $\text{Dgm}_k(\mathbb{K}; \varphi)$ defined in Eq. (3.A.2) is in $\mathcal{D}$. The following definition and stability result of the *bottleneck distance* are from Cohen-Steiner et al. (2007), and they will play important roles in the proof of Theorem 3.2.3.

Definition 3.A.1. Let $\mathbb{K}$ be a compact topological space. φ_1 and φ_2 are two continuous real-valued functions on $\mathbb{K}$ such that $\mathbb{K}$ is tame with respect to both φ_1 and φ_2 . The bottleneck distance between PDs $\text{Dgm}_k(\mathbb{K}; \varphi_1)$ and $\text{Dgm}_k(\mathbb{K}; \varphi_2)$ is defined as

$$W_\infty\left(\text{Dgm}_k(\mathbb{K}; \varphi_1), \text{Dgm}_k(\mathbb{K}; \varphi_2)\right) \stackrel{\text{def}}{=} \inf_{\gamma} \left(\sup \left\{ \|\xi - \gamma(\xi)\|_{l^\infty} \mid \xi \in \text{Dgm}_k(\mathbb{K}; \varphi_1) \right\} \right),$$

where γ ranges over bijections from $\text{Dgm}_k(\mathbb{K}; \varphi_1)$ to $\text{Dgm}_k(\mathbb{K}; \varphi_2)$, and

$$\|\xi\|_{l^\infty} \stackrel{\text{def}}{=} \max\{|\xi_1|, |\xi_2|\}, \quad \text{for all } \xi = (\xi_1, \xi_2)^\top \in \mathbb{R}^2. \quad (3.A.3)$$

Theorem 3.A.1. Let $\mathbb{K}$ be a compact and triangulable topological space. φ_1 and φ_2 are two continuous real-valued functions on $\mathbb{K}$ such that $\mathbb{K}$ is tame with respect to both φ_1 and φ_2 . Then, we have the bottleneck distance as follows

$$W_\infty\left(\text{Dgm}_k(\mathbb{K}; \varphi_1), \text{Dgm}_k(\mathbb{K}; \varphi_2)\right) \leq \sup_{x \in \mathbb{K}} |\varphi_1(x) - \varphi_2(x)|.$$

3.B Computation of SECT

Let $K \subset \mathbb{R}^d$ be a shape of interest. Suppose a finite set of points $\{x_i\}_{i=1}^I \subset K$ and a radius $r > 0$ are properly chosen such that

$$K_t^\nu = \{x \in K \mid x \cdot \nu \leq t - R\} \approx \bigcup_{i \in \mathfrak{I}_t^\nu} \overline{B(x_i, r)}, \quad \text{for all } t \in [0, T] \text{ and } \nu \in \mathbb{S}^{d-1}, \quad (3.B.1)$$

$$\text{where } \mathfrak{I}_t^\nu \stackrel{\text{def}}{=} \left\{ i \in \mathbb{N} \mid 1 \leq i \leq I \text{ and } x_i \cdot \nu \leq t - R \right\},$$

and $\overline{B(x_i, r)} := \{x \in \mathbb{R}^d : \|x - x_i\| \leq r\}$ denotes a closed ball centered at x_i with radius r . For example, when $d = 2$, centers x_i may be chosen as a subset of the grid points

$$\left\{ y_{j,j'} \stackrel{\text{def}}{=} (-R + j \cdot \delta, -R + j' \cdot \delta)^\top \right\}_{j,j'=1}^J$$

of the square $[-R, R]^2$ containing shape K , where $\delta = \frac{2R}{J}$ and radius $r = \delta$. Specifically,

$$K_t^\nu \approx \bigcup_{y_{j,j'} \in K_t^\nu} \overline{B(y_{j,j'}, \delta)} \quad \text{for all } t \in [0, T] \text{ and } \nu \in \mathbb{S}^{d-1}, \quad (3.B.2)$$

which is a special case of Eq. (3.B.1). The shape approximation in Eq. (3.B.2) is illustrated by Figures 3.5(a) and (b).

Čech complexes The Čech Complex determined by the point set $\{x_i\}_{i \in \mathcal{I}_t^\nu}$ and radius r in Eq. (3.B.1) is defined as the following simplicial complex

$$\check{C}_r(\{x_i\}_{i \in \mathcal{I}_t^\nu}) \stackrel{\text{def}}{=} \left\{ \text{conv}(\{x_i\}_{i \in s}) \mid s \in 2^{\mathcal{I}_t^\nu} \text{ and } \bigcap_{i \in s} \overline{B(x_i, r)} \neq \emptyset \right\},$$

where $\text{conv}(\{x_i\}_{i \in s})$ denotes the convex hull generated by points $\{x_i\}_{i \in s}$. The nerve theorem (see Chapter III of Edelsbrunner and Harer (2010)) indicates that the Čech Complex $\check{C}_r(\{x_i\}_{i \in \mathcal{I}_t^\nu})$ and the union $\bigcup_{i \in \mathcal{I}_t^\nu} \overline{B(x_i, r)}$ have the same homotopy type. Hence, they share the same Betti numbers, i.e.,

$$\beta_k(\check{C}_r(\{x_i\}_{i \in \mathcal{I}_t^\nu})) = \beta_k\left(\bigcup_{i \in \mathcal{I}_t^\nu} \overline{B(x_i, r)}\right), \quad \text{for all } k \in \mathbb{Z}.$$

Using the shape approximation in Eq. (3.B.1), we have the following approximation for ECC

$$\chi_t^\nu(K) \approx \sum_{k=0}^{d-1} (-1)^k \cdot \beta_k(\check{C}_r(\{x_i\}_{i \in \mathcal{I}_t^\nu})), \quad t \in [0, T]. \quad (3.B.3)$$

The method of computing the Betti numbers of simplicial complexes in Eq. (3.B.3) is standard and can be found in the literature (e.g., Chapter IV of Edelsbrunner and Harer (2010) and Section 3.1 of Niyogi et al. (2008)). Then, the SECT of K is estimated using Eq. (3.2.3). The smoothing effect of the integrals in Eq. (3.2.3) reduces the estimation error.

Computing the SECT for our proof-of-concept and simulation examples For the shape $K^{(1)}$ defined in Eq. (3.2.9) of Section 3.2.2, we estimate the SECT of $K^{(1)}$ using the aforementioned Čech complex approach with the following setup: $R = \frac{3}{2}$, $r = \frac{1}{5}$, and the point set $\{x_i\}_i$ is equal to the following collection

$$\left\{ \left(\frac{2}{5} + \cos t_j, \sin t_j \right) \mid t_j = \frac{\pi}{5} + \frac{j}{J} \cdot \frac{8\pi}{5} \right\}_{j=1}^J \cup \left\{ \left(-\frac{2}{5} + \cos t_j, \sin t_j \right) \mid \frac{6\pi}{5} + \frac{j}{J} \cdot \frac{8\pi}{5} \right\}_{j=1}^J, \quad (3.B.4)$$

where J is a sufficiently large integer. We implement $J = 100$ in our proof-of-concept example. Figure 3.5(c) illustrates the shape approximation using this setup. The SECT for other shapes in our proof-of-concept/simulation

examples is estimated in the similar way.

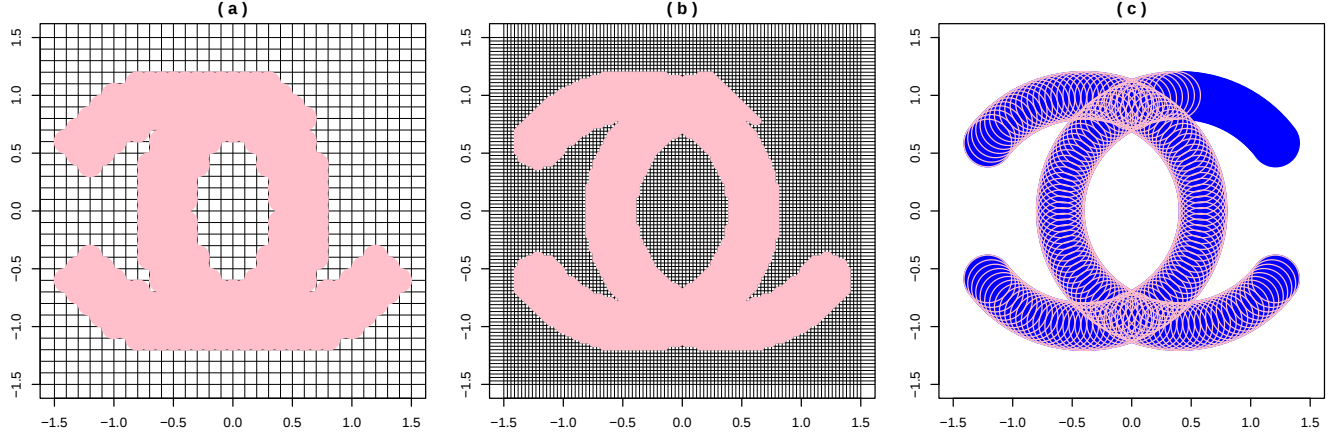


Figure 3.5: Illustrations of the shape approximation in Eq. (3.B.1). The shape K of interest herein is equal to the $K^{(1)}$ defined in Eq. (3.2.9) (see Figure 3.1(a) and the blue shape in the panel (c) herein). We set $R = \frac{3}{2}$, $t = 1$, and $\nu = (\sqrt{2}/2, \sqrt{2}/2)^\top$. Panels (a) and (b) specifically illustrate the approximation in Eq. (3.B.2) using grid points, and the pink shapes in the two panels present the union $\bigcup_{y_{j,j'} \in K_t^\nu} \overline{B(y_{j,j'}, \delta)}$ in Eq. (3.B.2); in panel (a), $J = 30$ (i.e., $\delta = 0.1$); in panel (b), $J = 100$ (i.e., $\delta = 0.03$). Panel (c) illustrates the approximation in Eq. (3.B.1) with centers $\{x_i\}_i$ being the points in Eq. (3.B.4) and $r = \frac{1}{5}$. Each pink circle in panel (c) presents a ball $\overline{B(x_i, r)}$.

3.C Proofs

Proof of Theorem 3.2.1 The following inclusion is straightforward

$$\{\text{Dgm}_k(K; \phi_\nu) \cap (-\infty, t) \times (t, \infty)\} \subset \{\xi \in \text{Dgm}_k(K; \phi_\nu) \mid \text{pers}(\xi) > 0\},$$

where the definitions of $\text{Dgm}_k(K; \phi_\nu)$ and $\text{pers}(\xi)$ are provided in Appendix 3.A. Together with the k-triangle lemma (see Edelsbrunner et al. (2000) and Cohen-Steiner et al. (2007)), this inclusion implies

$$\begin{aligned} \beta_k(K_t^\nu) &= \#\{\text{Dgm}_k(K; \phi_\nu) \cap (-\infty, t) \times (t, \infty)\} \\ &\leq \#\{\xi \in \text{Dgm}_k(K; \phi_\nu) \mid \text{pers}(\xi) > 0\} \leq \frac{M}{d}, \end{aligned}$$

for all $k \in \{1, \dots, d-1\}$, all $\nu \in \mathbb{S}^{d-1}$, and all t that are not HCPs in direction ν , where $\#\{\cdot\}$ counts the multiplicity of the multisets. Then, result (i) comes from the RCLL property of functions $\{\beta_k(K_t^\nu)\}_{t \in [0, T]}$. Eq. (3.2.2) implies

$$\begin{aligned} \sup_{\nu \in \mathbb{S}^{d-1}} \left(\sup_{0 \leq t \leq T} |\chi_t^\nu(K)| \right) &= \sup_{\nu \in \mathbb{S}^{d-1}} \left(\sup_{0 \leq t \leq T} \left| \sum_{k=0}^{d-1} (-1)^k \cdot \beta_k(K_t^\nu) \right| \right) \\ &\leq \sup_{\nu \in \mathbb{S}^{d-1}} \left[\sup_{0 \leq t \leq T} \left(d \cdot \sup_{k \in \{0, \dots, d-1\}} \beta_k(K_t^\nu) \right) \right] \leq M, \end{aligned}$$

which follows from result (i) and gives result (ii). $\square$

Proof of Theorem 3.2.2 Because of

$$\begin{aligned} \left\{ \int_0^t \chi_\tau^\nu(K) d\tau \right\}_{t \in [0, T]} &\in \{\text{all absolutely continuous functions on } [0, T]\} \\ &= \{x \in L^1([0, T]) : \text{the weak derivative } x' \text{ exists and } x' \in L^1([0, T])\} \\ &\stackrel{\text{def}}{=} W^{1,1}([0, T]) \end{aligned}$$

(see the Remark 8 after Proposition 8.3 in [Brezis \(2011\)](#) for details), the weak derivative of $\{\int_0^t \chi_\tau^\nu(K) d\tau\}_{t \in [0, T]}$ exists. This weak derivative is easily computed using the tameness of K , integration by parts, and the definition of weak derivatives. The proof of result (i) is completed. For the simplicity of notations, we denote

$$F(t) \stackrel{\text{def}}{=} \int_0^t \chi_\tau^\nu(K) d\tau - \frac{t}{T} \int_0^T \chi_\tau^\nu(K) d\tau, \quad \text{for } t \in [0, T].$$

Theorem 3.2.1 implies

$$|F(t)| \leq \int_0^T |\chi_\tau^\nu(K)| d\tau + \frac{t}{T} \int_0^T |\chi_\tau^\nu(K)| d\tau \leq 2TM, \quad \text{for } t \in [0, T].$$

Hence, $F \in L^p([0, T])$ for $p \in [1, \infty)$. Result (i) implies that the weak derivative of F exists and is $F'(t) = \chi_t^\nu(K) - \frac{1}{T} \int_0^T \chi_\tau^\nu(K) d\tau$. We have the boundedness

$$|F'(t)| \leq |\chi_t^\nu(K)| + \frac{1}{T} \int_0^T |\chi_\tau^\nu(K)| d\tau \leq 2M, \quad \text{for } t \in [0, T],$$

which implies $F' \in L^p([0, T])$ for $p \in [1, \infty)$. Furthermore, $F(0) = F(T) = 0$, together with the discussion above, implies $F \in W_0^{1,p}([0, T])$ for all $p \in [1, \infty)$ (see the Theorem 8.12 of [Brezis \(2011\)](#)). Theorem 8.8 and the Remark 8 after Proposition 8.3 in [Brezis \(2011\)](#) imply $W_0^{1,p}([0, T]) \subset \mathcal{B}$ for $p \in [1, \infty)$. Result (ii) follows. $\square$

The following lemmas are prepared for the proof of Theorem 3.2.3.

Lemma 3.C.1. *Suppose $K \in \mathcal{S}_{R,d}^M$. Then, we have the following estimate for all t that are neither HCPs in direction ν_1 nor HCPs in direction ν_2 .*

$$\begin{aligned} \Upsilon_k(t; \nu_1, \nu_2) &\stackrel{\text{def}}{=} |\beta_k(K_t^{\nu_1}) - \beta_k(K_t^{\nu_2})| \\ &\leq \# \left\{ x \in \text{Dgm}_k(K; \phi_{\nu_1}) \mid x \neq \gamma^*(x) \text{ and } \underline{(x, \gamma^*(x))} \cap \partial((-\infty, t) \times (t, \infty)) \neq \emptyset \right\}, \end{aligned} \tag{3.C.1}$$

where $\underline{(x, \gamma^*(x))}$ denotes the straight line segment connecting points x and $\gamma^*(x)$ in $\mathbb{R}^2$, γ^* is any optimal bijection

such that

$$W_\infty\left(\text{Dgm}_k(K; \phi_{\nu_1}), \text{Dgm}_k(K; \phi_{\nu_2})\right) = \sup \left\{ \|\xi - \gamma^*(\xi)\|_{l^\infty} \mid \xi \in \text{Dgm}_k(\mathbb{K}; \phi_{\nu_1}) \right\} \quad (3.C.2)$$

(see Definition 3.A.1, and $\|\cdot\|_{l^\infty}$ is defined in Eq. (3.A.3)), and “#” counts the corresponding multiplicity.

Remark 3.C.1. Because $(\mathcal{D}, W_\infty)$ is a geodesic space, the optimal bijection γ^* does exist (see Proposition 1 of Turner (2013) and its proof therein).

Proof. Since t is not an HCP, neither $\text{Dgm}_k(K; \phi_{\nu_1})$ nor $\text{Dgm}_k(K; \phi_{\nu_2})$ has a point on the boundary $\partial((-\infty, t) \times (t, \infty))$. If $\beta_k(K_t^{\nu_1}) = \beta_k(K_t^{\nu_2})$, Eq. (3.C.1) is true. Otherwise, without loss of generality, we assume $\beta_k(K_t^{\nu_1}) > \beta_k(K_t^{\nu_2})$. Notice

$$\beta_k(K_t^{\nu_i}) = \# \left\{ \text{Dgm}_k(K; \phi_{\nu_i}) \cap ((-\infty, t) \times (t, \infty)) \right\}, \quad \text{for } i \in \{1, 2\}.$$

Let γ^* be any optimal bijection, then there should be at least $\beta_k(K_t^{\nu_1}) - \beta_k(K_t^{\nu_2})$ straight line segments $\underline{(x, \gamma^*(x))}$ crossing $\partial((-\infty, t) \times (t, \infty))$; otherwise, γ^* is not bijective. Hence,

$$\begin{aligned} & \beta_k(K_t^{\nu_1}) - \beta_k(K_t^{\nu_2}) \\ & \leq \# \left\{ x \in \text{Dgm}_k(K; \phi_{\nu_1}) \mid x \neq \gamma^*(x) \text{ and } \underline{(x, \gamma^*(x))} \cap \partial((-\infty, t) \times (t, \infty)) \neq \emptyset \right\}, \end{aligned}$$

and Eq. (3.C.1) follows. □

Lemma 3.C.2. Suppose $K \in \mathcal{S}_{R,d}^M$. Except for finitely many t , we have

$$\Upsilon_k(t; \nu_1, \nu_2) \leq \frac{2M}{d} \cdot \mathbf{1}_{\mathcal{T}_k}, \quad \text{where}$$

$$\begin{aligned} \mathcal{T}_k \stackrel{\text{def}}{=} \left\{ t \in [0, T] \text{ not an HCP in direction } \nu_1 \text{ or } \nu_2 \mid \text{there exists } x \in \text{Dgm}_k(K; \phi_{\nu_1}) \text{ such that } x \neq \gamma^*(x) \right. \\ \left. \text{and } \underline{(x, \gamma^*(x))} \cap \partial((-\infty, t) \times (t, \infty)) \neq \emptyset \right\}, \end{aligned}$$

and $\gamma^* : \text{Dgm}_k(K; \phi_{\nu_1}) \rightarrow \text{Dgm}_k(K; \phi_{\nu_2})$ is any optimal bijection satisfying Eq. (3.C.2).

Proof. Theorem 3.2.1 implies

$$\Upsilon_k(t; \nu_1, \nu_2) = |\beta_k(K_t^{\nu_1}) - \beta_k(K_t^{\nu_2})| \leq 2M/d.$$

Furthermore, the inequality in Eq. (3.C.1) indicates that $\Upsilon_k(t; \nu_1, \nu_2) = 0$ if $t \notin \mathcal{T}_k$, except for finitely many HCPs in directions ν_1 and ν_2 . Then the desired estimate follows. □

Proof of Theorem 3.2.3 The definition of Euler characteristic (see Eq. (3.2.2)) and Lemma 3.C.2 imply the following: for $p \in [1, \infty)$, we have

$$\begin{aligned}
& \int_0^T \left| \left\{ \chi_\tau^{\nu_1}(K) - \chi_\tau^{\nu_2}(K) \right\} \right|^p d\tau \\
&= \int_0^T \left| \sum_{k=0}^{d-1} (-1)^k \cdot \left(\beta_k(K_\tau^{\nu_1}) - \beta_k(K_\tau^{\nu_2}) \right) \right|^p d\tau \\
&\leq \int_0^T \left(\sum_{k=0}^{d-1} \Upsilon_k(\tau; \nu_1, \nu_2) \right)^p d\tau \\
&\leq d^{(p-1)} \cdot \sum_{k=0}^{d-1} \int_0^T \left(\Upsilon_k(\tau; \nu_1, \nu_2) \right)^p d\tau \\
&\leq \frac{(2M)^p}{d} \cdot \sum_{k=0}^{d-1} \int_{\mathcal{T}_k} d\tau \\
&\leq \frac{(2M)^p}{d} \cdot \sum_{k=0}^{d-1} \left(\sum_{\xi \in \text{Dgm}_k(K; \phi_{\nu_1})} 2 \cdot \|\xi - \gamma^*(\xi)\|_{l^\infty} \right),
\end{aligned} \tag{3.C.3}$$

where the last inequality follows from the definition of $\mathcal{T}_k$. Since $\|\xi - \gamma^*(\xi)\|_{l^\infty}$ can be positive only if $\text{pers}(\xi) > 0$ or $\text{pers}(\gamma^*(\xi)) > 0$, there are at most N terms $\|\xi - \gamma^*(\xi)\|_{l^\infty} > 0$, where condition (3.2.4) implies

$$N \stackrel{\text{def}}{=} \sum_{i=1}^2 \#\{\xi \in \text{Dgm}_k(K; \phi_{\nu_i}) \mid \text{pers}(\xi) > 0\} \leq 2M/d.$$

Therefore, the inequality in Eq. (3.C.3) implies

$$\begin{aligned}
\int_0^T \left| \left\{ \chi_\tau^{\nu_1}(K) - \chi_\tau^{\nu_2}(K) \right\} \right|^p d\tau &\leq \frac{2 \cdot (2M)^{(p+1)}}{d} \cdot \sup \left\{ \|\xi - \gamma^*(\xi)\|_{l^\infty} \mid \xi \in \text{Dym}(\mathbb{K}; \phi_{\nu_1}) \right\} \\
&= \frac{2 \cdot (2M)^{(p+1)}}{d} \cdot W_\infty \left(\text{Dgm}_k(K; \phi_{\nu_1}), \text{Dgm}_k(K; \phi_{\nu_2}) \right).
\end{aligned}$$

Then, Theorem 3.A.1 implies

$$\int_0^T \left| \left\{ \chi_\tau^{\nu_1}(K) - \chi_\tau^{\nu_2}(K) \right\} \right|^p d\tau \leq \frac{2 \cdot (2M)^{(p+1)}}{d} \cdot \sup_{x \in K} |x \cdot (\nu_1 - \nu_2)|.$$

Additionally, $|x \cdot (\nu_1 - \nu_2)| \leq \|x\| \cdot \|\nu_1 - \nu_2\|$ and $K \subset \overline{B(0, R)}$ provide

$$\int_0^T \left| \left\{ \chi_\tau^{\nu_1}(K) - \chi_\tau^{\nu_2}(K) \right\} \right|^p d\tau \leq \frac{2 \cdot R \cdot (2M)^{(p+1)}}{d} \cdot \|\nu_1 - \nu_2\|. \tag{3.C.4}$$

Define the constant $C_{M,R,d}^*$ as follows

$$C_{M,R,d}^* \stackrel{\text{def}}{=} \sqrt{\frac{16M^3R}{d} + \frac{32M^3R}{d} + \frac{64M^4R}{d^2}}.$$

Setting $p = 2$, Eq. (3.C.4) implies the following

$$\left(\int_0^T \left| \left\{ \chi_\tau^{\nu_1}(K) - \chi_\tau^{\nu_2}(K) \right\} \right|^2 d\tau \right)^{1/2} \leq \sqrt{\frac{16M^3R}{d} \cdot \|\nu_1 - \nu_2\|} \leq C_{M,R,d}^* \cdot \sqrt{\|\nu_1 - \nu_2\|},$$

which is the inequality in Eq. (3.2.6).

The definition of $SECT(K)$, together with Eq. (3.C.4), implies

$$\begin{aligned} & \left\| SECT(K)(\nu_1) - SECT(K)(\nu_2) \right\|_{\mathcal{H}}^2 \\ &= \int_0^T \left| \frac{d}{dt} SECT(K)(\nu_1; t) - \frac{d}{dt} SECT(K)(\nu_2; t) \right|^2 dt \\ &= \int_0^T \left| \left(\chi_t^{\nu_1}(K) - \chi_t^{\nu_2}(K) \right) - \frac{1}{T} \int_0^T \left(\chi_\tau^{\nu_1}(K) - \chi_\tau^{\nu_2}(K) \right) d\tau \right|^2 dt \\ &\leq \int_0^T \left(\left| \chi_t^{\nu_1}(K) - \chi_t^{\nu_2}(K) \right| + \frac{1}{T} \int_0^T \left| \chi_\tau^{\nu_1}(K) - \chi_\tau^{\nu_2}(K) \right| d\tau \right)^2 dt \\ &\leq 2 \int_0^T \left| \chi_t^{\nu_1}(K) - \chi_t^{\nu_2}(K) \right|^2 dt + \frac{2}{T} \left(\int_0^T \left| \chi_\tau^{\nu_1}(K) - \chi_\tau^{\nu_2}(K) \right| d\tau \right)^2 \\ &\leq \frac{32M^3R}{d} \cdot \|\nu_1 - \nu_2\| + \frac{64M^4R}{d^2} \cdot \|\nu_1 - \nu_2\|^2, \end{aligned}$$

where the last inequality above comes from Eq. (3.C.4). Then, we have

$$\begin{aligned} & \left\| SECT(K)(\nu_1) - SECT(K)(\nu_2) \right\|_{\mathcal{H}} \\ &\leq \sqrt{\frac{32M^3R}{d} \cdot \|\nu_1 - \nu_2\| + \frac{64M^4R}{d^2} \cdot \|\nu_1 - \nu_2\|^2} \\ &\leq C_{M,R,d}^* \cdot \sqrt{\|\nu_1 - \nu_2\| + \|\nu_1 - \nu_2\|^2}. \end{aligned}$$

The proof of the second inequality in result (i) is completed.

The law of cosines and Taylor's expansion indicates

$$\begin{aligned} \frac{\|\nu_1 - \nu_2\|}{d_{\mathbb{S}^{d-1}}(\nu_1, \nu_2)} &= \sqrt{2 \cdot \frac{1 - \cos(d_{\mathbb{S}^{d-1}}(\nu_1, \nu_2))}{\{d_{\mathbb{S}^{d-1}}(\nu_1, \nu_2)\}^2}} \\ &= \sqrt{2 \cdot \left[\sum_{n=1}^{\infty} \frac{(-1)^{n+1}}{(2n)!} \cdot \{d_{\mathbb{S}^{d-1}}(\nu_1, \nu_2)\}^{2n-2} \right]} = O(1). \end{aligned}$$

Then, result (ii) comes from the following

$$\begin{aligned} & \frac{\|SECT(K)(\nu_1) - SECT(K)(\nu_2)\|_{\mathcal{H}}}{\sqrt{d_{\mathbb{S}^{d-1}}(\nu_1, \nu_2)}} \\ & \leq C_{M,R,d}^* \cdot \sqrt{\frac{\|\nu_1 - \nu_2\|}{d_{\mathbb{S}^{d-1}}(\nu_1, \nu_2)} + \frac{\|\nu_1 - \nu_2\|^2}{d_{\mathbb{S}^{d-1}}(\nu_1, \nu_2)}} = O(1). \end{aligned} \quad (3.C.5)$$

We consider the following inequality for all $\nu_1, \nu_2 \in \mathbb{S}^{d-1}$ and $t_1, t_2 \in [0, T]$

$$\begin{aligned} & |SECT(K)(\nu_1; t_1) - SECT(K)(\nu_2; t_2)| \\ & \leq |SECT(K)(\nu_1; t_1) - SECT(K)(\nu_1; t_2)| \\ & \quad + |SECT(K)(\nu_1; t_2) - SECT(K)(\nu_2; t_2)| \\ & \stackrel{\text{def}}{=} I + II. \end{aligned} \quad (3.C.6)$$

From the definition of $\|\cdot\|_{C^{0, \frac{1}{2}}([0, T])}$ and Eq. (3.1.1), we have

$$\begin{aligned} & \sup_{t_1, t_2 \in [0, T], t_1 \neq t_2} \frac{|SECT(K)(\nu_1; t_1) - SECT(K)(\nu_1; t_2)|}{|t_1 - t_2|^{1/2}} \\ & \leq \|SECT(K)(\nu_1)\|_{C^{0, \frac{1}{2}}([0, T])} \\ & \leq \tilde{C}_T \|SECT(K)(\nu_1)\|_{\mathcal{H}} \\ & \leq \tilde{C}_T \|SECT(K)\|_{C(\mathbb{S}^{d-1}; \mathcal{H})}, \end{aligned}$$

which implies $I \leq \tilde{C}_T \|SECT(K)\|_{C(\mathbb{S}^{d-1}; \mathcal{H})} \cdot |t_1 - t_2|^{1/2}$ for all $t_1, t_2 \in [0, T]$. Applying Eq. (3.1.1) again, we have

$$\begin{aligned} II & \leq \|SECT(K)(\nu_1) - SECT(K)(\nu_2)\|_{\mathcal{B}} \\ & \leq \tilde{C}_T \|SECT(K)(\nu_1) - SECT(K)(\nu_2)\|_{\mathcal{H}}. \end{aligned}$$

Then, the inequality in Eq. (3.2.7) follows from Eq. (3.C.6) and result (ii). With an argument similar to Eq. (3.C.5), the function $(\nu, t) \mapsto SECT(K)(\nu; t)$ belongs to $C^{0, \frac{1}{2}}(\mathbb{S}^{d-1} \times [0, T]; \mathbb{R})$. The proof of Theorem 3.2.3 is completed. $\square$

The proof of Theorem 3.2.4 needs the following concepts, which can be found in [Andradas et al. \(2012\)](#).

Definition 3.C.1. (*Constructible sets*) (i) A locally closed set is a subset of a topological space that is the intersection of an open and a closed subset. (ii) A constructible set is a finite union of the aforementioned locally closed sets.

Proof of Theorem 3.2.4 Since $K \in \mathcal{S}_{R,d}^M$ is a triangulable subset of $\mathbb{R}^d$, the shape K is homeomorphic to a polygon $|S| = \bigcup_{s \in S} s \subset \mathbb{R}^d$, where S is a finite simplicial complex and each s is a simplex. We may assume $K = |S|$

without loss of generality. Since $s = s \cap \mathbb{R}^d$, each simplex s is a locally closed set; hence, K is constructible (see Definition 3.C.1). Then, Corollary 6 of Ghrist et al. (2018) implies the desired result. $\square$

Proof of Theorem 3.3.1 The triangle inequalities and symmetry of ρ follow from that of the metric of $C(\mathbb{S}^{d-1}; \mathcal{H})$. Equation $\rho(K_1, K_2) = 0$ indicates $\|PECT(K_1)(\nu) - PECT(K_2)(\nu)\|_{\mathcal{H}_{BM}} = 0$ for all $\nu \in \mathbb{S}^{d-1}$. Evans (2010) (Theorem 5 of Chapter 5.6) implies $\|PECT(K_1)(\nu) - PECT(K_2)(\nu)\|_{\mathcal{B}} = 0$ for all $\nu \in \mathbb{S}^{d-1}$. Then, we have $\int_0^t \chi_\tau^\nu(K_1) d\tau = \int_0^t \chi_\tau^\nu(K_2) d\tau$ for all $t \in [0, T]$ and $\nu \in \mathbb{S}^{d-1}$; hence, $SECT(K_1) = SECT(K_2)$. Then the invertibility in Theorem 3.2.4 implies $K_1 = K_2$. Therefore, ρ is a distance, and the proof of result (i) is completed.

The proof of result (ii) is motivated by the following chain of maps for any fixed $\nu \in \mathbb{S}^{d-1}$ and $t \in [0, T]$.

$$\begin{aligned} \mathcal{S}_{R,d}^M &\xrightarrow{PECT} C(\mathbb{S}^{d-1}; \mathcal{H}_{BM}) \xrightarrow{\text{projection}} \mathcal{H}_{BM}, \text{ which is embedded into } \mathcal{B} \xrightarrow{\text{projection}} \mathbb{R}, \\ K &\mapsto \{PECT(K)(\nu')\}_{\nu' \in \mathbb{S}^{d-1}} \mapsto \{PECT(K)(\nu; t')\}_{t' \in [0, T]} \mapsto PECT(K)(\nu; t) = \int_0^t \chi_\tau^\nu(K) d\tau, \end{aligned}$$

where all spaces above are metric spaces and equipped with their Borel algebras. We notice the following facts:

- mapping $PECT : \mathcal{S}_{R,d}^M \rightarrow C(\mathbb{S}^{d-1}; \mathcal{H}_{BM})$ is isometric;
- projection $C(\mathbb{S}^{d-1}; \mathcal{H}_{BM}) \rightarrow \mathcal{H}_{BM}$, $\{F(\nu')\}_{\nu' \in \mathbb{S}^{d-1}} \mapsto F(\nu)$ is continuous for each fixed direction ν ;
- applying Evans (2010) (Theorem 5 of Chapter 5.6) again, the embedding $\mathcal{H}_{BM} \rightarrow \mathcal{B}$, $F(\nu) \mapsto F(\nu)$ is continuous;
- projection $\mathcal{B} \rightarrow \mathbb{R}$, $\{x(t')\}_{t' \in [0, T]} \mapsto x(t)$ is continuous.

Therefore, $\mathcal{S}_{R,d}^M \rightarrow \mathbb{R}$, $K \mapsto PECT(K)(\nu, t)$ is continuous, hence measurable. Because $\chi_{(\cdot)}^\nu(K)$, for each $K \in \mathcal{S}_{R,d}^M$, is a RCLL step function with finitely many discontinuities,

$$\chi_t^\nu(K) = \lim_{n \rightarrow \infty} \left[\frac{1}{\delta_n} \{PECT(K)(\nu; t + \delta_n) - PECT(K)(\nu; t)\} \right],$$

for all $K \in \mathcal{S}_{R,d}^M$, where $\lim_{n \rightarrow \infty} \delta_n = 0$ and $\delta_n > 0$. The measurability of $PECT(\cdot)(\nu; t + \delta_n)$ and $PECT(\cdot)(\nu; t)$ implies that $\chi_t^\nu(\cdot) : \mathcal{S}_{R,d}^M \rightarrow \mathbb{R}$, $K \mapsto \chi_t^\nu(K)$ is measurable, for any fixed ν and t . The proof of result (ii) is completed. $\square$

Proof of Theorem 3.3.3 For each fixed direction $\nu \in \mathbb{S}^{d-1}$, Theorem 3.3.2 indicates that the mapping $SECT(\cdot)(\nu) : K \mapsto SECT(K)(\nu)$ is an $\mathcal{H}$ -valued measurable function defined on the probability space $(\mathcal{S}_{R,d}^M, \mathcal{F}, \mathbb{P})$. We first show the Bochner $\mathbb{P}$ -integrability of $SECT(\cdot)(\nu)$ (see Section 5 in Chapter V of Yosida et al. (1965) for the definition of Bochner $\mathbb{P}$ -integrability), and the Bochner integral of $SECT(\cdot)(\nu)$ will be fundamental to our proof of Theorem 3.3.3. Lemma 1.3 of Da Prato and Zabczyk (2014) indicates that $SECT(\cdot)(\nu)$ is strongly $\mathcal{F}$ -measurable

(see Section 4 in Chapter V of [Yosida et al. \(1965\)](#) for the definition of strong $\mathcal{F}$ -measurability). Then, Assumption 2 indicates that the Bochner integral

$$m_\nu^* \stackrel{\text{def}}{=} \int_{\mathcal{S}_{R,d}^M} SECT(K)(\nu) \mathbb{P}(dK)$$

is Bochner $\mathbb{P}$ -integrable and $m_\nu^* \in \mathcal{H}$ (see [Yosida et al. \(1965\)](#), Section 5 of Chapter V, Theorem 1 therein particularly). The Corollary 2 in Section 5 of Chapter V of [Yosida et al. \(1965\)](#), together with that $\mathcal{H}$ is the RKHS generated by the kernel $\kappa(\cdot, \cdot)$ defined in Eq. (3.2.5), implies

$$\begin{aligned} m_\nu^*(t) &= \langle \kappa(t, \cdot), m_\nu^* \rangle \\ &= \int_{\mathcal{S}_{R,d}^M} \langle \kappa(t, \cdot), SECT(K)(\nu) \rangle \mathbb{P}(dK) \\ &= \int_{\mathcal{S}_{R,d}^M} SECT(K)(\nu; t) \mathbb{P}(dK) \\ &= \mathbb{E} \{ SECT(\cdot)(\nu; t) \} = m_\nu(t), \quad \text{for all } t \in [0, T]. \end{aligned}$$

Therefore, $m_\nu = m_\nu^* \in \mathcal{H}$. The proof of result (i) is completed.

To prove result (ii), we first show the product measurability of the following map for each fixed direction $\nu \in \mathbb{S}^{d-1}$

$$\begin{aligned} \left(\mathcal{S}_{R,d}^M \times [0, T], \mathcal{F} \times \mathcal{B}([0, T]) \right) &\rightarrow (\mathbb{R}, \mathcal{B}(\mathbb{R})), \\ (K, t) &\mapsto SECT(K)(\nu; t), \end{aligned} \tag{3.C.7}$$

where $\mathcal{F} \times \mathcal{B}([0, T])$ denotes the product σ -algebra generated by $\mathcal{F}$ and $\mathcal{B}([0, T])$. Define the filtration $\{\mathcal{F}_t\}_{t \in [0, T]}$ by $\mathcal{F}_t = \sigma(\{SECT(\cdot)(\nu; t') \mid t' \in [0, t]\})$ for $t \in [0, T]$. Because the paths of $SECT(\cdot)(\nu)$ are in $\mathcal{H}$, these paths are continuous (see Eq. (3.1.2)). Proposition 1.13 of [Karatzas and Shreve \(2012\)](#) implies that the stochastic process $SECT(\cdot)(\nu)$ is progressively measurable with respect to the filtration $\{\mathcal{F}_t\}_{t \in [0, T]}$. Then, the mapping in Eq. (3.C.7) is measurable with respect to the product σ -algebra $\mathcal{F} \times \mathcal{B}([0, T])$ (see [Karatzas and Shreve \(2012\)](#), particularly Definitions 1.6 and 1.11, also the paragraph right after Definition 1.11 therein). Eq. (3.3.4) implies

$$\int_0^T \int_{\mathcal{S}_{R,d}^M} |SECT(K)(\nu; t)|^2 \mathbb{P}(dK) dt \leq T \cdot \tilde{C}_T^2 \cdot \mathbb{E} \|SECT(\cdot)(\nu)\|_{\mathcal{H}}^2 < \infty,$$

where the double integral is well-defined because of the product measurability of the mapping in Eq. (3.C.7) and the Fubini's theorem. Then, the proof of result (ii) is completed.

For any $\nu_1, \nu_2 \in \mathbb{S}^{d-1}$, the proof of result (i) implies the following Bochner integral representation

$$\begin{aligned}
\|m_{\nu_1} - m_{\nu_2}\|_{\mathcal{H}} &= \left\| \int_{\mathcal{S}_{R,d}^M} SECT(K)(\nu_1) - SECT(K)(\nu_2) \mathbb{P}(dK) \right\|_{\mathcal{H}} \\
&\stackrel{(1)}{\leq} \int_{\mathcal{S}_{R,d}^M} \left\| SECT(K)(\nu_1) - SECT(K)(\nu_2) \right\|_{\mathcal{H}} \mathbb{P}(dK) \\
&\stackrel{(2)}{\leq} C_{M,R,d}^* \cdot \sqrt{\|\nu_1 - \nu_2\| + \|\nu_1 - \nu_2\|^2},
\end{aligned}$$

where the inequality (1) follows from the Corollary 1 in Section 5 of Chapter V of [Yosida et al. \(1965\)](#), and the inequality (2) follows from Theorem [3.2.3](#) (i). With the argument in Eq. [\(3.C.5\)](#), the proof of result (iii) is completed. $\square$

Bibliography

- R. A. Adams and J. J. Fournier. *Sobolev spaces. 2nd ed.* Amsterdam: Academic Press, 2003.
- G. K. Aguirre, E. Zarahn, and M. D’Esposito. Empirical analyses of bold fmri statistics. *NeuroImage*, 5(3):179–197, 1997.
- A. Alexanderian. A brief note on the karhunen–loève expansion. *arXiv preprint arXiv:1509.07526*, 2015.
- C. Andradas, L. Bröcker, and J. M. Ruiz. *Constructible sets in real geometry*, volume 33. Springer Science & Business Media, 2012.
- F. G. Ashby. *Statistical analysis of fMRI data*. MIT press, 2019.
- D. M. Barch, G. C. Burgess, M. P. Harms, S. E. Petersen, B. L. Schlaggar, M. Corbetta, M. F. Glasser, S. Curtiss, S. Dixit, C. Feldt, et al. Function in the human connectome: task-fmri and individual differences in behavior. *Neuroimage*, 80:169–189, 2013.
- M. Belkin and P. Niyogi. Laplacian eigenmaps for dimensionality reduction and data representation. *Neural Computation*, 15(6):1373–1396, 2003.
- A. Berline and C. Thomas-Agnan. *Reproducing kernel Hilbert spaces in probability and statistics*. Springer Science & Business Media, 2011.
- E. Boissard, T. Le Gouic, J.-M. Loubes, et al. Distribution’s template estimate with wasserstein metrics. *Bernoulli*, 21(2):740–759, 2015.
- D. M. Boyer, Y. Lipman, E. S. Clair, J. Puente, B. A. Patel, T. Funkhouser, J. Jernvall, and I. Daubechies. Algorithms to automatically quantify the geometric similarity of anatomical surfaces. *Proceedings of the National Academy of Sciences*, 108(45):18221–18226, 2011.
- H. Brezis. *Functional analysis, Sobolev spaces and partial differential equations*, volume 2. Springer, 2011.
- R. L. Buckner, F. M. Krienen, A. Castellanos, J. C. Diaz, and B. T. Yeo. The organization of the human cerebellum estimated by intrinsic functional connectivity. *Journal of neurophysiology*, 106(5):2322–2345, 2011.

- Y. Chen, Z. Lin, and H.-G. Müller. Wasserstein regression. *Journal of the American Statistical Association*, pages 1–14, 2021.
- S.-s. Chern. A simple intrinsic proof of the gauss-bonnet formula for closed riemannian manifolds. *Annals of mathematics*, pages 747–752, 1944.
- J. M. Cisler, K. Bush, and J. S. Steele. A comparison of statistical methods for detecting context-modulated functional connectivity in fmri. *Neuroimage*, 84:1042–1052, 2014.
- J. Cohen. A coefficient of agreement for nominal scales. *Educational and psychological measurement*, 20(1):37–46, 1960.
- D. Cohen-Steiner, H. Edelsbrunner, and J. Harer. Stability of persistence diagrams. *Discrete & computational geometry*, 37(1):103–120, 2007.
- D. Cohen-Steiner, H. Edelsbrunner, J. Harer, and Y. Mileyko. Lipschitz functions have l_p -stable persistence. *Foundations of computational mathematics*, 10(2):127–139, 2010.
- J. B. Conway. *Functions of One Complex Variable I*. Springer Verlag, 1973.
- L. Crawford, A. Monod, A. X. Chen, S. Mukherjee, and R. Rabadán. Predicting clinical outcomes in glioblastoma: an application of topological and functional data analysis. *Journal of the American Statistical Association*, 115(531):1139–1150, 2020.
- I. Cribben and M. Fiecas. Functional connectivity analyses for fmri data. *Handbook of neuroimaging data analysis*, 369, 2016.
- J. Curry, S. Mukherjee, and K. Turner. How many directions determine a shape and other sufficiency results for two topological transforms. *arXiv preprint arXiv:1805.09782*, 2018.
- G. Da Prato and J. Zabczyk. *Stochastic equations in infinite dimensions*. Cambridge university press, 2014.
- D. A. Darling and P. Erdős. A limit theorem for the maximum of normalized sums of independent random variables. *Duke Mathematical Journal*, 23(1):143–155, 1956.
- A. P. Dempster, N. M. Laird, and D. B. Rubin. Maximum likelihood from incomplete data via the em algorithm. *Journal of the Royal Statistical Society: Series B (Methodological)*, 39(1):1–38, 1977.
- M. P. Do Carmo. *Differential geometry of curves and surfaces: revised and updated second edition*. New York : Dover Publications, INC., 2016.
- T. Duchamp, W. Stuetzle, et al. Extremal properties of principal curves in the plane. *The Annals of Statistics*, 24(4):1511–1520, 1996.

- J. Duchon. Splines minimizing rotation-invariant semi-norms in sobolev spaces. In *Constructive theory of functions of several variables*, pages 85–100. Springer, 1977.
- E. Dudek and K. Holly. Nonlinear orthogonal projection. In *Annales Polonici Mathematici*, volume 59, pages 1–31. Instytut Matematyczny Polskiej Akademii Nauk, 1994.
- D. B. Dunson and N. Wu. Inferring manifolds from noisy data using gaussian processes. *arXiv preprint arXiv:2110.07478*, 2021.
- H. Edelsbrunner and J. Harer. *Computational topology: an introduction*. American Mathematical Soc., 2010.
- H. Edelsbrunner, D. Letscher, and A. Zomorodian. Topological persistence and simplification. In *Proceedings 41st annual symposium on foundations of computer science*, pages 454–463. IEEE, 2000.
- A. Eloyan and S. K. Ghosh. Smooth density estimation with moment constraints using mixture distributions. *Journal of nonparametric statistics*, 23(2):513–531, 2011.
- K. Enomoto and M. Okura. The total squared curvature of curves and approximation by piecewise circular curves. *Results in Mathematics*, 64(1-2):215–228, 2013.
- L. C. Evans. *Partial differential equations*, volume 19. American Mathematical Soc., 2010.
- J. Fan. Test of significance based on wavelet thresholding and neyman’s truncation. *Journal of the American Statistical Association*, 91(434):674–688, 1996.
- B. T. Fasy, F. Lecci, A. Rinaldo, L. Wasserman, S. Balakrishnan, and A. Singh. Confidence sets for persistence diagrams. *The Annals of Statistics*, 42(6):2301–2339, 2014.
- M. Fréchet. Les éléments aléatoires de nature quelconque dans un espace distancié. In *Annales de l’institut Henri Poincaré*, volume 10, pages 215–310, 1948.
- K. Friston, C. Frith, P. Liddle, and R. Frackowiak. Functional connectivity: the principal-component analysis of large (pet) data sets. *Journal of Cerebral Blood Flow & Metabolism*, 13(1):5–14, 1993.
- K. J. Friston, P. Jezzard, and R. Turner. Analysis of functional mri time-series. *Human brain mapping*, 1(2):153–171, 1994.
- T. Gao, S. Z. Kovalsky, and I. Daubechies. Gaussian process landmarking on manifolds. *SIAM Journal on Mathematics of Data Science*, 1(1):208–236, 2019.
- S. Gerber and R. Whitaker. Regularization-free principal curve estimation. *The Journal of Machine Learning Research*, 14(1):1285–1302, 2013.

- R. Ghrist, R. Levanger, and H. Mai. Persistent homology and euler integral transforms. *arXiv preprint arXiv:1804.04740*, 2018.
- M. Hairer. An introduction to stochastic pdes. *arXiv preprint arXiv:0907.4178*, 2009.
- T. Hastie. Principal curves and surfaces. Technical report, Stanford University, California, Laboratory for Computational Statistics, 1984.
- T. Hastie and W. Stuetzle. Principal curves. *Journal of the American Statistical Association*, 84(406):502–516, 1989.
- A. Hatcher. *Algebraic topology*. Cambridge ; New York : Cambridge University Press, 2002.
- S. Hauberg. Principal curves on riemannian manifolds. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 38(9):1915–1921, 2015.
- R. M. Hutchison, T. Womelsdorf, E. A. Allen, P. A. Bandettini, V. D. Calhoun, M. Corbetta, S. Della Penna, J. H. Duyn, G. H. Glover, J. Gonzalez-Castillo, et al. Dynamic functional connectivity: promise, issues, and interpretations. *Neuroimage*, 80:360–378, 2013.
- S. E. Joel, B. S. Caffo, P. C. van Zijl, and J. J. Pekar. On the relationship between seed-based and ica-based measures of functional connectivity. *Magnetic Resonance in Medicine*, 66(3):644–657, 2011.
- I. T. Jolliffe. *Principal component analysis. 2nd ed.* New York: Springer, 2002.
- J. Jumper, R. Evans, A. Pritzel, T. Green, M. Figurnov, O. Ronneberger, K. Tunyasuvunakool, R. Bates, A. Žídek, A. Potapenko, et al. Highly accurate protein structure prediction with alphafold. *Nature*, 596(7873):583–589, 2021.
- O. Kallenberg. *Foundations of modern probability*, volume 2. Springer, 2021.
- I. Karatzas and S. Shreve. *Brownian motion and stochastic calculus*, volume 113. Springer Science & Business Media, 2012.
- B. Kégl, A. Krzyzak, T. Linder, and K. Zeger. Learning and design of principal curves. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 22(3):281–297, 2000.
- J.-H. Kim, J. Lee, and H.-S. Oh. Spherical principal curves. *arXiv preprint arXiv:2003.02578*, 2020.
- S. Kirov and D. Slepčev. Multiple penalized principal curves: Analysis and computation. *Journal of Mathematical Imaging and Vision*, 59(2):234–256, 2017.
- H. Kirveslahti and S. Mukherjee. Representing fields without correspondences: the lifted euler characteristic transform. *arXiv preprint arXiv:2111.04788*, 2021.

- A. Klenke. *Probability theory: a comprehensive course*. Springer Science & Business Media, 2013.
- R. Koenker and I. Mizera. Penalized triograms: Total variation regularization for bivariate smoothing. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)*, 66(1):145–163, 2004.
- H. Lee, M. K. Chung, H. Kang, B.-N. Kim, and D. S. Lee. Discriminative persistent homology of brain networks. In *2011 IEEE international symposium on biomedical imaging: from nano to macro*, pages 841–844. IEEE, 2011.
- S.-P. Lee, T. Q. Duong, G. Yang, C. Iadecola, and S.-G. Kim. Relative changes of cerebral arterial and venous blood volumes during increased cerebral blood flow: implications for bold fmri. *Magnetic Resonance in Medicine: An Official Journal of the International Society for Magnetic Resonance in Medicine*, 45(5):791–800, 2001.
- M. Lifshits. Lectures on gaussian processes. In *Lectures on Gaussian Processes*, pages 1–117. Springer, 2012.
- M. A. Lindquist et al. The statistical analysis of fmri data. *Statistical science*, 23(4):439–464, 2008.
- B. G. Lindsay. The geometry of mixture likelihoods: A general theory. *The Annals of Statistics*, pages 86–94, 1983.
- M. J. Lowe, M. Dzemidzic, J. T. Lurito, V. P. Mathews, and M. D. Phillips. Correlations in low-frequency bold fluctuations reflect cortico-cortical connections. *Neuroimage*, 12(5):582–587, 2000.
- L. K. Lynch, K.-H. Lu, H. Wen, Y. Zhang, A. J. Saykin, and Z. Liu. Task-evoked functional connectivity does not explain functional connectivity differences between rest and task conditions. *Human brain mapping*, 39(12):4939–4948, 2018.
- A. F. Mejia, M. B. Nebel, A. D. Barber, A. S. Choe, J. J. Pekar, B. S. Caffo, and M. A. Lindquist. Improved estimation of subject-level functional connectivity using full and partial correlation with empirical bayes shrinkage. *NeuroImage*, 172:478–491, 2018.
- K. Meng and A. Eloyan. Population-level task-evoked functional connectivity. *arXiv preprint arXiv:2102.12039*, 2021a.
- K. Meng and A. Eloyan. Principal manifold estimation via model complexity selection. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)*, 83(2):369–394, 2021b.
- K. Meng, L. Crawford, and A. Eloyan. Randomness and statistical inference of shapes via the smooth euler characteristic transform. *arXiv preprint arXiv:2204.12699*, 2022a.
- K. Meng, L. Crawford, and A. Eloyan. Supplement to “randomness and statistical inference of shapes via the smooth euler characteristic transform”. 2022b.
- F. M. Miezin, L. Maccotta, J. Ollinger, S. Petersen, and R. Buckner. Characterizing the hemodynamic response: effects of presentation rate, sampling procedure, and the possibility of ordering brain activity based on relative timing. *Neuroimage*, 11(6):735–759, 2000.

- Y. Mileyko, S. Mukherjee, and J. Harer. Probability measures on the space of persistence diagrams. *Inverse Problems*, 27(12):124007, 2011.
- C. Moon, Q. Li, and G. Xiao. Predicting survival outcomes using topological features of tumor pathology images. *arXiv e-prints*, pages arXiv–2012, 2020.
- K. Müller, G. Lohmann, V. Bosch, and D. Y. Von Cramon. On multivariate spectral analysis of fmri time series. *NeuroImage*, 14(2):347–356, 2001.
- J. R. Munkres. *Elements of algebraic topology*. CRC press, 2018.
- R. Nieuwenhuys, J. Hans, and C. Nicholson. *The central nervous system of vertebrates*. Springer, 2014.
- P. Niyogi, S. Smale, and S. Weinberger. Finding the homology of submanifolds with high confidence from random samples. *Discrete & Computational Geometry*, 39(1-3):419–441, 2008.
- A. Petersen and H.-G. Müller. Fréchet regression for random objects with euclidean predictors. *The Annals of Statistics*, 47(2):691–719, 2019.
- R Core Team. *R: A Language and Environment for Statistical Computing*. R Foundation for Statistical Computing, Vienna, Austria, 2019. URL <https://www.R-project.org>.
- C. E. Rasmussen and C. K. Williams. *Gaussian processes for machine learning*, volume 2. MIT press Cambridge, MA, 2006.
- M. Reed. *Methods of modern mathematical physics: Functional analysis*. Elsevier, 2012.
- J. Rissman, A. Gazzaley, and M. D’Esposito. Measuring functional connectivity during distinct stages of a cognitive task. *Neuroimage*, 23(2):752–763, 2004.
- A. Robinson and K. Turner. Hypothesis testing for topological data analysis. *Journal of Applied and Computational Topology*, 1(2):241–261, 2017.
- S. T. Roweis and L. K. Saul. Nonlinear dimensionality reduction by locally linear embedding. *Science*, 290(5500):2323–2326, 2000.
- W. Rudin. Functional analysis. 1991. *Internat. Ser. Pure Appl. Math*, 1991.
- V. Seguy and M. Cuturi. Principal geodesic analysis for probability measures under the optimal transport metric. *Advances in Neural Information Processing Systems*, 28:3312–3320, 2015.
- A. J. Smola, S. Mika, B. Schölkopf, and R. C. Williamson. Regularized principal manifolds. *Journal of Machine Learning Research*, 1(Jun):179–209, 2001.

- E. Somasundaram, A. Litzler, R. Wadhwa, S. Owen, and J. Scott. Persistent homology of tumor ct scans is associated with survival in lung cancer. *Medical physics*, 48(11):7043–7051, 2021.
- J. B. Tenenbaum, V. De Silva, and J. C. Langford. A global geometric framework for nonlinear dimensionality reduction. *Science*, 290(5500):2319–2323, 2000.
- R. Tibshirani. Principal curves revisited. *Statistics and Computing*, 2(4):183–190, 1992.
- L. Tong, R.-W. Liu, V. C. Soon, and Y.-F. Huang. Indeterminacy and identifiability of blind identification. *IEEE Transactions on circuits and systems*, 38(5):499–509, 1991.
- K. Turner. Means and medians of sets of persistence diagrams. *arXiv preprint arXiv:1307.8300*, 2013.
- K. Turner, Y. Mileyko, S. Mukherjee, and J. Harer. Fréchet means for distributions of persistence diagrams. *Discrete & Computational Geometry*, 52(1):44–70, 2014a.
- K. Turner, S. Mukherjee, and D. M. Boyer. Persistent homology transform for modeling shapes and surfaces. *Information and Inference: A Journal of the IMA*, 3(4):310–344, 2014b.
- N. Tzourio-Mazoyer, B. Landeau, D. Papathanassiou, F. Crivello, O. Etard, N. Delcroix, B. Mazoyer, and M. Joliot. Automated anatomical labeling of activations in spm using a macroscopic anatomical parcellation of the mni mri single-subject brain. *Neuroimage*, 15(1):273–289, 2002.
- O. Vipond, J. A. Bull, P. S. Macklin, U. Tillmann, C. W. Pugh, H. M. Byrne, and H. A. Harrington. Multiparameter persistent homology landscapes identify immune cell spatial patterns in tumors. *Proceedings of the National Academy of Sciences*, 118(41), 2021.
- G. Wahba. *Spline models for observational data*, volume 59. SIAM, 1990.
- B. Wang, T. Sudijono, H. Kirveslahti, T. Gao, D. M. Boyer, S. Mukherjee, and L. Crawford. A statistical pipeline for identifying physical features that differentiate classes of 3d shapes. *The Annals of Applied Statistics*, 15(2): 638–661, 2021.
- R. Warnick, M. Guindani, E. Erhardt, E. Allen, V. Calhoun, and M. Vannucci. A bayesian approach for estimating dynamic functional network connectivity in fmri data. *Journal of the American Statistical Association*, 113(521): 134–151, 2018.
- J. Yang, I. Anishchenko, H. Park, Z. Peng, S. Ovchinnikov, and D. Baker. Improved protein structure prediction using predicted interresidue orientations. *Proceedings of the National Academy of Sciences*, 117(3):1496–1503, 2020.

- B. T. Yeo, F. M. Krienen, J. Sepulcre, M. R. Sabuncu, D. Lashkari, M. Hollinshead, J. L. Roffman, J. W. Smoller, L. Zöllei, J. R. Polimeni, et al. The organization of the human cerebral cortex estimated by intrinsic functional connectivity. *Journal of neurophysiology*, 2011.
- K. Yosida et al. Functional analysis. 1965.
- C. Yue, V. Zipunnikov, P.-L. Bazin, D. Pham, D. Reich, C. Crainiceanu, and B. Caffo. Parameterization of white matter manifold-like structures using principal surfaces. *Journal of the American Statistical Association*, 111(515):1050–1060, 2016.
- T. Zhang, F. Li, L. Beckes, and J. A. Coan. A semi-parametric model of the hemodynamic response for multi-subject fmri data. *NeuroImage*, 75:136–145, 2013.