
SCIENTIFIC REPORTS

Frontiers of Functional Data Analysis: Challenges and Opportunities in the Era of AI

19 Aug 2024–13 Sep 2024

Organizing Committee

Alexander Aue

University of California at Davis

Ying Chen

National University of Singapore

Zhenhua Lin

National University of Singapore

Qiwei Yao

London School of Economics

CONTENTS PAGE

		Page
Karthik Bharath University of Nottingham, UK	Rolling-Without-Slipping Models for Manifold-Valued Functional Data	4
Jinyuan Chang Southwestern University of Finance and Economics, China	On the Modelling and Prediction of High-Dimensional Functional Time Series	8
Maria Grith Erasmus University Rotterdam, Netherlands	Group Factors in Single Stock Options	9
Giles Hooker University of Pennsylvania, USA	Functional Data Analysis as Nonparametric ODEs	10
Tailen Hsing University of Michigan, USA	Estimating the Spectral Density of a Function-Valued Spatial Process	14
Jian Huang The Hong Kong Polytechnic University, Hong Kong SAR	Continuous Normalizing Flows for Learning Probability Distributions and Short Course on Statistical Deep Learning	16
Hannah Lai National University of Singapore, Singapore	Neural Tangent Kernel in Implied Volatility Forecasting: A Nonlinear Functional Autoregression Approach	18
Jialiang Li National University of Singapore, Singapore	Robust Model Averaging Prediction	21
Hans-Georg Müller University of California, Davis, USA	Modelling Distribution-Valued Random Trajectories with Optimal Transports	22
Byeong Uk Park Seoul National University, S.Korea	High-Dimensional Hilbert-Schmidt Linear Regression for Hilbert Manifold Variables	26
Xinghao Qiao The University of Hong Kong, Hong Kong SAR	Large-Scale Multiple Testing of Cross-Covariance Functions with Applications to Functional Network Models	28
Eftychia Solea Queen Mary University of London, UK	Robust Inverse Regression for Multivariate Elliptical Functional Data	31
Jane-Ling Wang University of California, Davis, USA	FDA in the age of AI	33
Wanjie Wang National University of Singapore, Singapore	Recovery of Time Labels for Noisy Dynamic Data	34
Fang Yao Peking University, China	FPCA with Diverging Dimensions: From Sparse to Dense Designs	35

	Page
Zhigang Yao National University of Singapore, Singapore	39
Jin-Ting Zhang National University of Singapore, Singapore	40

ROLLED MODELS FOR MANIFOLD-VALUED FUNCTIONAL DATA

KARTHIK BHARATH

Classification AMS 2020: 62R30; 62M20; 62H12

Keywords: Fréchet mean; Gaussian process; Parallel transport

1. INTRODUCTION

Imagine that a curve γ on the unit sphere $\mathbb{S}^2$ drawn in wet ink is rolled along a plane without slipping or twisting so as to trace out a curve $\gamma^\downarrow$ in the tangent space of the initial point $\gamma(0)$, identified with $\mathbb{R}^2$. The geometric operation engenders a local isometry between the two curves so that they determine each other uniquely upto isometries, and the operation may be extended to arbitrary d -dimensional connected complete manifolds M . Mathematically, in the intrinsic picture, the Euclidean curve $\gamma^\downarrow : [0, 1] \rightarrow T_{\gamma(0)}M$ on the tangent space of the starting point $\gamma(0)$ is known as the *unrolling* of γ , and is determined by the initial value differential equation

$$\dot{\gamma}^\downarrow(t) = P_{0 \leftarrow t}^\gamma \dot{\gamma}(t), \quad \gamma^\downarrow(0) = 0,$$

where $\dot{\gamma}(t) = \frac{d}{dt}\gamma(t) \in T_{\gamma(t)}M$, and $P_{0 \leftarrow t}^\gamma : T_{\gamma(t)}M \rightarrow T_{\gamma(0)}M$ is the parallel transport map along the curve γ , a linear isometry. Choice of coordinates in a tangent space is arbitrary, a frame is hence needed to represent a tangent vector in standard coordinates of $\mathbb{R}^d$. Given an orthonormal frame $U : T_{\gamma(0)}M \rightarrow \mathbb{R}^d$, $U\gamma^\downarrow(t)$ is then a curve in $\mathbb{R}^d$. In fact, $\gamma^\downarrow$ may be defined on the tangent space T_bM of an arbitrary point $b \in M$ outside the cut locus of $\gamma(0)$ by modifying the differential equation as

$$(1) \quad \dot{\gamma}^\downarrow(t) = P_{0 \leftarrow 1}^c P_{0 \leftarrow t}^\gamma \dot{\gamma}(t), \quad \gamma^\downarrow(0) = \exp_b^{-1}(\gamma(0)),$$

where $\exp^{-1}$ is the inverse of Riemannian exponential map $\exp : TM \rightarrow M$, and c is the geodesic between b and $\gamma(0)$. Absence of slipping is characterised by use of the parallel transport along γ ; twisting is relevant when $\mathbb{S}^2$ is viewed as an embedded submanifold of $\mathbb{R}^3$, where the parallel transport in the above equation is replaced by a curve in $SE(3)$, the isometry group of $\mathbb{R}^3$ [1].

Operationally, from (1), for a fixed point $b \in M$ equipped with a frame U for its tangent space, we can define four maps to: (i) unroll a curve γ in M to obtain a curve $\gamma^\downarrow$ in $\mathbb{R}^d$; (ii) perform the reverse operation of *rolling* an $\mathbb{R}^d$ -valued curve $\gamma^\downarrow$ to obtain γ on M ; (iii) use the exponential map at $\gamma(t)$, to *wrap* a curve z in $\mathbb{R}^d$ with respect to γ in M by parallel transporting along curves c and γ the deviation $Uz - \gamma^\downarrow$ from the mean; (iv) *unwrap* a curve x on M with respect γ on M by reversing the wrapping operation. Under some conditions, the wrapping and unwrapping maps are also inverses of each other. Rolling/unrolling operations have been used in statistics for curve-fitting using splines, first on $\mathbb{S}^2$ [2] and more recently on general manifolds [3].

The goal of this work is to use the four maps to: define generative statistical models for functional data assuming values in M by pushing forward under the rolling and

wrapping maps a parametric stochastic process model $\{\mathbb{P}_\theta, \theta \in \Theta\}$ for random curves in $\mathbb{R}^d$; given discretely observed M -valued data $\{x_i(t_j)\}$ on a common time grid, estimate θ using unrolling and unwrapping maps. In particular, we will focus on the case where $\{\mathbb{P}_\theta, \theta \in \Theta\}$ corresponds to Gaussian measures parametrized by a mean and positive definite covariance function. Figure 1 provides an illustration.

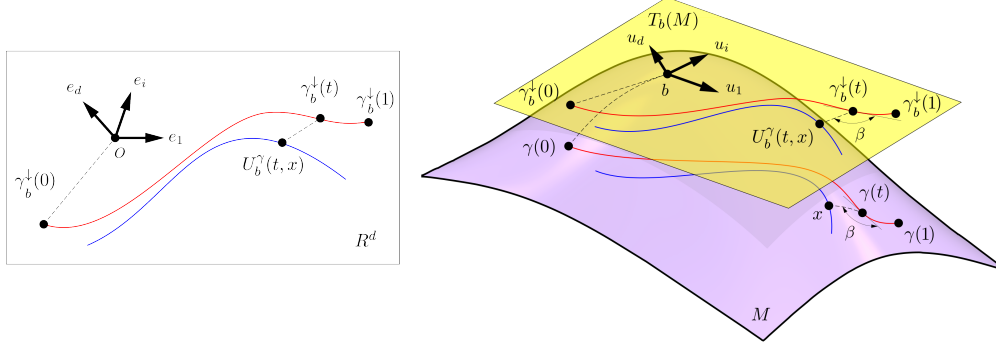


FIGURE 1. Realisation from a Gaussian process in $\mathbb{R}^d$ is mapped to a random curve on M . Red is the mean of the Gaussian procecess, with respect to which the rolling is performed, and blue is the realization from the Gaussian proceess that deviates from the mean curve. The line segment connecting points between the blue and red curves at arbitrary t , and the corresponding angle denoted β , indicate distances and angles preserved by the (un)rolling.

2. THEORETICAL RESULTS

2.1. Fréchet mean and rolled mean. For a random curve $x : [0, 1] \rightarrow M$ the *population Fréchet mean curve* is defined as the Fréchet mean of x , defined pointwise in t as the minimizer of $s \mapsto E\{\rho^2(x(t), s)\}$, where ρ is the intrinsic distance on M ; its sample version is defined by taking expectation with respect to the empirical measure on a sample of curves. The Fréchet mean curve coincides with the rolling of the mean of the $\mathbb{R}^d$ -valued process, not necessarily Gaussian, under some conditions.

Theorem 2.1. *For every t , the rolled mean is the Fréchet mean of $x(t)$ if any of the following conditions are true:*

- (i) *Every point in M has an empty cut locus;*
- (ii) *M is a symmetric manifold, the distribution of $x(t)$ has even symmetry about $\gamma(t)$, and has a unique Fréchet mean that lies outside the cut locus of $\gamma(t)$.*

Condition (ii) concerns the interplay between notions of symmetry of a manifold and that of a probability measure on it. A function $g : M \rightarrow \mathbb{R}$ on a symmetric manifold M is said to be symmetric if $g(p) = g(\sigma_p(p))$ for every $p \in M$, where $\sigma_p : M \rightarrow M$ is a geodesic-reversing isometry. A distribution ν on M is said to possess *even symmetry* about $p \in M$ if $\nu = (\exp_p)_\# \lambda$, the pushforward under the exponential map at p of a mean-zero distribution λ on $T_p M$ with Lebesgue density f , when $T_p M$ is identified with $\mathbb{R}^d$, satisfies $f(v) = f(-v)$ for every $v \in T_p M$.

2.2. Rolled Gaussian process on M . Consider a Gaussian process $t \mapsto z(t) \in \mathbb{R}^d$ with mean function $t \mapsto m(t) \in \mathbb{R}^d$ and covariance kernel $(s, t) \mapsto k(s, t) \in \text{Sym}_{>0}(d)$, where $\text{Sym}_{>0}(d)$ is the cone of positive definite matrices within the vector space of $d \times d$ real symmetric matrices. The wrapping map may be used, with respect to the rolled mean, to transform z to a stochastic process on M , which we refer to as a *rolled Gaussian process*.

Definition 2.2. Let $z \sim GP(m, k)$, and choose $b \in M$ and frame U of $T_b M$. With $\tilde{m}_b = Um$ as a curve in $T_b M$, let $\gamma := \tilde{m}_b^\uparrow$ be the rolling of $\tilde{m}$. The process $x = y_b^{\uparrow\gamma}$ obtained by wrapping of $y_b := Uz$ with respect to γ is a rolled Gaussian process, denoted $x \sim \mathcal{RGP}(m, k; b, U)$.

The point b and a frame U for $T_b M$ are arbitrary, but inconsequential for modelling.

Proposition 2.3. Starting from point $b \in M$ with frame U for $T_b M$, let $x \sim \mathcal{RGP}(m, K; b, U)$. If one starts instead from $b' \in M$ with basis U' for $T_{b'} M$, then there exists unique m' and K' such that the rolled Gaussian process $x' \sim \mathcal{RGP}(m', K'; b', U')$ is equal in distribution to x .

For practical purposes, it is convenient to consider a parametric model for the Gaussian process z with respect to a particular basis. Let $\{\phi_s : [0, 1] \rightarrow \mathbb{R}\}$ be a B-spline basis [4]. Let the mean $m(t) = M_w \phi(t)$, and assume a separable covariance $K(t, t') = \phi(t)^\top V_w \phi(t') U_w$, where $\phi(t) = \{\phi_1(t), \dots, \phi_k(t)\} \in \mathbb{R}^k$ is a vector and $M_w \in \mathbb{R}^{d \times k}$, $U_w \in \text{Sym}_{>0}(d)$ and $V_w \in \text{Sym}_{>0}(k)$ are matrices that parameterise the model. The convenience of this particular choice is that the curve z can be written

$$(2) \quad z(t) = \sum_{s=1}^k w_s \phi_s(t),$$

where $W = (w_1, \dots, w_k) \sim \mathcal{MN}(M_w, U_w, V_w)$, the matrix normal distribution with mean matrix, M_w , row covariance, U_w , and column covariance, V_w .

2.3. Estimation. If $Z \in \mathbb{R}^{d \times r}$ is obtained by observing z in (2) at times $t_1, \dots, t_r$ then $Z \sim \mathcal{MN}(M_w \Phi, U_w, \Phi^\top V_w \Phi)$, where $\Phi = \{\phi(t_1), \dots, \phi(t_r)\} \in \mathbb{R}^{k \times r}$; the distribution of X , the corresponding discretisation of the rolled Gaussian process x is denoted as $X \sim \mathcal{RMN}(M_w, U_w, V_w; b, U)$, and represents the model for the discretely observed sample of curves $\{x_i(t_j)\}$.

Exploiting the relationship between the rolled mean and the Fréchet mean curves in Theorem 2.1, the estimator $\hat{M}_w$ of the mean parameter M_w is defined as follows. Let $H(\hat{\Gamma}) \in \mathbb{R}^{d \times r}$ be the unrolling of the discretised sample Fréchet mean curve onto $T_b M$ followed by a transformation to standard coordinates using the frame U . Define $\hat{M}_w = H(\hat{\Gamma})\Phi^-$, where Φ^- is the right inverse of Φ .

Theorem 2.4. Let $x \sim \mathcal{RGP}(m, k; b, U)$. Assume that the Fréchet mean curve of x exists and is unique, and suppose that the sample Fréchet mean curve converges in probability to it, as $n \rightarrow \infty$, in the C^1 topology. Then, under any of the conditions in Theorem 2.1, as $n \rightarrow \infty$, $\hat{M}_w$ converges in probability to M_w .

The C^1 topology for convergence is needed to ensure that the sequence of parallel transport maps along the sample Fréchet mean curve converge to their limit along the the population Fréchet mean curve. Estimators of covariances U_w and V_w are defined using $\hat{M}_w$ [5], but it is unclear if they are consistent.

3. ROBOTICS APPLICATION WITH CURVES ON $SO(3)$

The data are time-indexed orientations of the end-effector of a Franka robot arm as it was guided $n = 60$ times to perform a task to deposit the contents of a dustpan into a bin at $r = 100$ time points. 3D orientations are represented by elements of the rotation group $SO(3)$, which under the unsigned unit quaternion representation, can be identified with $\mathbb{S}^3$ modulo the antipodal map. Fixing the sign of each data point then identifies $SO(3)$ with a hemisphere of $\mathbb{S}^3$.

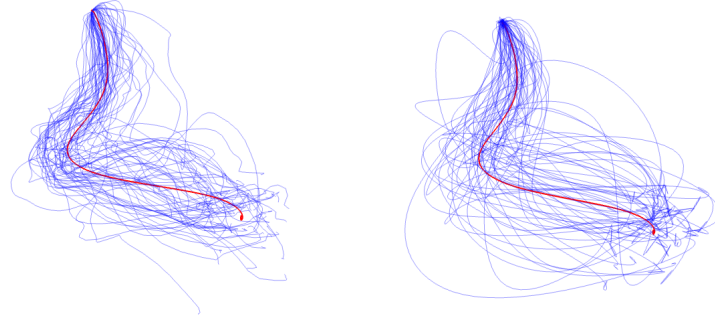


FIGURE 2. Left: Unwrapped $SO(3)$ curves (blue), and unrolled fitted mean (red); Right: simulations from the fitted Gaussian process model;

Left panel of Figure 2 shows the unwrapped curves (blue) in $\mathbb{R}^3$ and the unrolled mean, $H(\hat{\Gamma})$ (red) based on $\hat{M}_w$. The curves have a common starting point at $t = 0$, shown near the top-left in this plot; variability seems to increase with t , especially following a kink point that corresponds to the dustpan being turned to empty its contents. As a visual appraisal of the fitted model, right panel of Figure 2 shows $n = 60$ realisations from $\mathcal{RMN}(\hat{M}_w, \hat{U}_w, \hat{V}_w; b, U)$, using the same unwrapping coordinates and projection as in left panel of Figure 2. These simulated curves are smoother than the real data, which is a consequence of the basis used, but they have similar heteroscedastic variation.

REFERENCES

- [1] R. W. Sharpe (2000). *Differential geometry: Cartan's generalization of Klein's Erlangen program*, Springer.
- [2] P. E. Jupp and J. T. Kent (1987). Fitting smooth paths to spherical data. *Journal of the Royal Statistical Society Series C: Applied Statistics*, 36, 34-36.
- [3] K. R. Kim, I. L. Dryden, H. Le, and K. E. Severn (2021). Smoothing splines on Riemannian manifolds, with applications to 3D shape space *Journal of the Royal Statistical Society Series B: Statistical Methodology*, 83, 108-132.
- [4] C. DeBoor (2002). *A practical guide to splines*, Springer.
- [5] P. Dutilleul (1999). The MLE algorithm for the matrix normal distribution. *Journal of statistical computation and simulation*, 64, 105-123.

SCHOOL OF MATHEMATICAL SCIENCES, UNIVERSITY OF NOTTINGHAM
Email address: karthik.bharath@nottingham.ac.uk

ON THE MODELLING AND PREDICTION OF HIGH-DIMENSIONAL FUNCTIONAL TIME SERIES

JINYUAN CHANG

Keywords: Dimension reduction; Eigenanalysis; Functional thresholding; Hilbert–Schmidt norm; Permutation; Segmentation transformation

We propose a two-step procedure to model and predict high-dimensional functional time series, where the number of function-valued time series p is large in relation to the length of time series n . Our first step performs an eigenanalysis of a positive definite matrix, which leads to a one-to-one linear transformation for the original high-dimensional functional time series, and the transformed curve series can be segmented into several groups such that any two subseries from any two different groups are uncorrelated both contemporaneously and serially. Consequently in our second step those groups are handled separately without the information loss on the overall linear dynamic structure. The second step is devoted to establishing a finite-dimensional dynamical structure for all the transformed functional time series within each group. Furthermore the finite-dimensional structure is represented by that of a vector time series. Modelling and forecasting for the original high-dimensional functional time series are realized via those for the vector time series in all the groups. We investigate the theoretical properties of our proposed methods, and illustrate the finite-sample performance through both extensive simulation and two real datasets.

<https://arxiv.org/abs/2406.00700>

REFERENCES

- [1] Chang, J., Guo, B. and Yao, Q. Principal component analysis for second-order stationary vector time series. *The Annals of Statistics*, 46, 2094–2124, 2018.
- [2] Fang, Q., Guo, S. and Qiao, X. Finite sample theory for high-dimensional functional/scalar time series with applications. *Electronic Journal of Statistics*, 16, 527–591, 2022.
- [3] Golub, G.H. and Van Loan, C.F. Matrix Computations. *Johns Hopkins Studies in the Mathematical Sciences*, fourth edn, Johns Hopkins University Press, Baltimore, MD., 1996.
- [4] Guo, S. and Qiao, X. On consistency and sparsity for high-dimensional functional time series with application to autoregressions. *Bernoulli*, 29, 451–472, 2023.

SOUTHWESTERN UNIVERSITY OF FINANCE AND ECONOMICS
Email address: changjinyuan@swufe.edu.cn

GROUP FACTORS IN SINGLE STOCK OPTIONS

MARIA GRITH

We consider a linear functional panel model in which the response variables are functions, and the regressors are scalar common factors. The functional response introduces a third dimension in the panel alongside the cross-section and time dimensions. We assume that the coefficient functions are characterized by a multilevel latent group pattern of heterogeneity. For discretely observed functions, we use penalized-sieve estimation. We estimate the group memberships and functional coefficients using an estimator that minimizes a least squares criterion with respect to all possible groupings of the crosssectional response functions when the number of groups is known. We provide conditions under which our estimators are consistent as the observations in the three dimensions of the panel tend to infinity, and we develop inference methods. Additionally, we show that using a functional setup improves clustering accuracy and convergence rates of the group-specific functional coefficients. The gains are enhanced if membership is factor-invariant. We also propose a generalized information criterion for estimating the number of groups when it is unknown. We evaluate the performance of our method in a simulation with both densely and sparsely observed functional responses. Finally, we apply our approach to identify grouped patterns of unobserved heterogeneity in a panel of single-stock options implied volatility surfaces.

ERASMUS UNIVERSITY ROTTERDAM
Email address: grith@ese.eur.nl

AN UNDERSTANDING OF PRINCIPAL DIFFERENTIAL ANALYSIS

GILES HOOKER AND EDWARD GUNNING

Classification AMS 2020: 62G08, 62H12, 62H25, 62M09

Keywords: principal differential analysis, Gaussian process, dimension reduction, ordinary differential equation

1. INTRODUCTION

Classically, functional data analysis focusses on the study of data comprised of univariate functions measured on some continuous domain: $X_1(t), \dots, X_n(t)$, although this conception has been frequently extended into a more general class of *object data* [3]. One of the distinguishing features of this framework is access to derivatives $D^k X(t)$; presenting both questions about selecting between derivatives as covariates [1, 2] as well as relationships between derivatives.

Principal Differential Analysis (PDA) was proposed in [4] as a means of modeling these relationships. In the classical PDA formulation, an m -th order ODE of the form:

$$(1.1) \quad D^m X(t) = \beta_0(t)X(t) + \beta_1(t)DX(t) + \dots + \beta_{m-1}(t)D^{m-1}X(t) + \eta(t)$$

is proposed as a model for data. This takes the form of a concurrent linear model described in [5] in which $D^m(X)$ is a response that depends on lower-order derivatives through time-varying functions $\beta(t)$. However, (1.1) also takes the form of a time-varying linear ordinary differential equation (ODE), and PDA leverages this framework in a number of ways.

In particular, PDA is proposed as providing two quantities:

(1) Data reduction: solutions to (1.1) can be expressed as

$$(1.2) \quad X(t) = \Phi(t, 0)x_0 + \int \Phi(t, s)\eta(s)ds$$

in which $\Phi(s, t)$ is an $(m-1) \times (m-1)$ *transition* matrix at each (s, t) and x_0 is the vector $(X(0), DX(0), \dots, D^{m-1}X(0))$. Here we regard $\Phi(t, 0)$ as providing a basis expansion in which to represent $X(t)$, providing a data-reduction method akin to functional principal components analysis.

(2) Representation of behaviour. For time-invariant ODE's in which the $\beta(t)$ are constant with $\eta(t) = 0$. Solutions to (1.1) can be expressed in terms of

$$(1.3) \quad X(t) = \sum_{k=1}^{m-1} c_k e^{b_k t} (\cos(d_k t) + i \sin(d_k t))$$

in which (b_k, c_k, d_k) are obtained from the eigen decomposition of a matrix representation of the dynamics. Under this framework, the instantaneous behaviour of $X(t)$ can be interpreted by considering the decomposition in (1.3) at the current values of β .

We revisit the PDA model, re-interpreting it as a description of a data-generating process in which $\eta(t)$ is a random error process specific to each observations. This view yields a number of consequences, most specifically a bias in the classical PDA estimates which we correct with an iterative procedure.

2. CONSEQUENCES OF A GENERATIVE MODEL

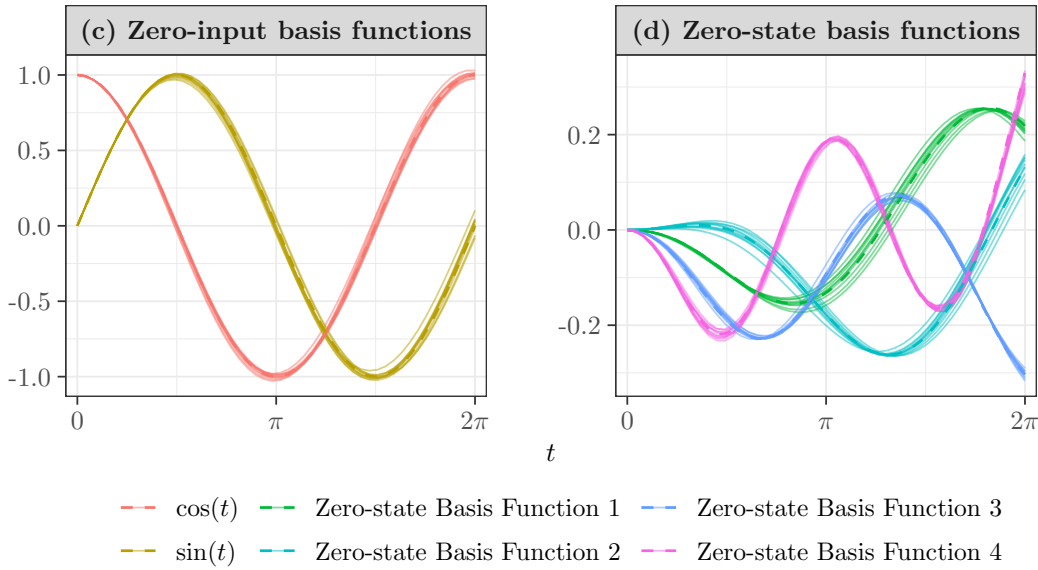
This presentation regards (1.1) as a generating model, we assume that $\eta(t) \sim (0, \Sigma)$ is a mean zero random error process with covariance $\Sigma(s, t)$. We first observe that this framework augments the dimension reduction approach in [4] with a second term derived from (1.2):

$$\text{cov}(\tilde{\mathbf{x}}(s), \tilde{\mathbf{x}}(t)) = \Phi(s, 0)\Sigma_0\Phi(t, 0)^\top + \int_0^s \int_0^t \Phi(s, u)\Sigma(u, v)\Phi(t, v)^\top dv du$$

yielding a new decomposition. We illustrate this below based on simulating data from simple harmonic motion forced by Gaussian process noise:

$$D^2X(t) = -X(t) + \eta(t)$$

illustrated below in which the left hand plot provides simulation estimates of variation associated with initial conditions $\mathbf{x}_0$ and the right-hand plot variance associated with the random process $\eta(\cdot)$.



A second consequence involves biases in the estimated $\hat{\beta}(t)$. The original PDA formulation minimizes the integrated sum of squared errors from (1.1):

$$\sum \int (D^m X_i(t) - \beta_0 X_i(t) - \dots - \beta_{m-1}(t) D^{m-1} X_i(t))^2 dt$$

which can be solved by a linear regression at each time point t . Writing $Z(t)$ as the matrix containing rows $(X(t), DX(t), \dots, D^{m-1}X(t))$ we obtain an estimate

$$(2.1) \quad \hat{\beta}(t) = (Z(t)^T Z(t))^{-1} Z(t)^T D^{m-1} X(t)$$

although [4] represented each $\beta_j(t)$ via a basis expansion allowing the use of smoothing penalties.

In classical linear regression (2.1) substituting $D^{m-1}X(t) = Z(t)\beta(t) + \eta(t)$ results in unbiased estimates. Here, however, we observe from (1.2) that $EZ(t)\eta(t) = \int_{s=0}^t \Phi(s, t)_{m-1, \cdot} \Sigma(s, t) ds$ and an approximation to bias in $\hat{\beta}$:

$$E\hat{\beta}(t) = \beta(t) + E(Z(t)^T Z(T))^{-1} EZ(t) \int_{s=0}^t \Phi(s, t)_{m-1, \cdot} \Sigma(s, t) ds.$$

Within the bias term, both $\Phi(s, t)$ and Σ require estimates for $\beta(t)$, resulting in the following iterative

- (1) **Initialize** Begin with ordinary least squares (OLS) estimates for the parameters $\beta_0(t), \dots, \beta_{m-1}(t)$ by minimizing the ISSE and obtain residuals $\hat{\eta}(t)$ and covariance

$$\hat{\Sigma}(s, t) = \frac{1}{n - m} \sum \hat{\eta}_i(s) \hat{\eta}_i(t)$$

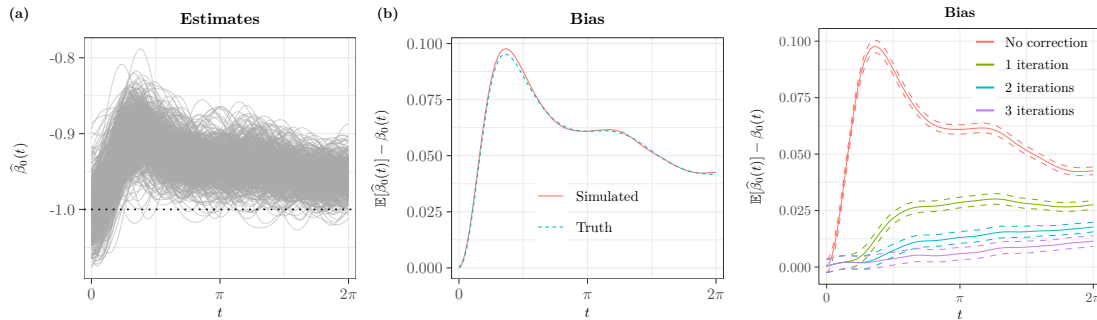
along with the transition matrices $\Phi(s, t)$.

- (2) **Iterative Correction:** Apply the bias-reduction formula iteratively:

$$(2.2) \quad \hat{\beta}_j^{(k+1)}(t) = \hat{\beta}_j^{(k)}(t) - E(Z(t)^T Z(T))^{-1} EZ(t) \int_{s=0}^t \Phi(s, t)_{m-1, \cdot} \Sigma(s, t) ds$$

After each iteration, the residuals $\hat{\eta}_i^{(k)}(t)$ are updated along with $\hat{\Sigma}$ and Φ until convergence.

We illustrate this with the same simple harmonic motion example below in which we first provide an estimate of the bias for the coefficient of $X(t)$ and estimates after three steps of bias reduction



3. CONCLUSION

This work reframes PDA as a statistical model that accounts for both deterministic ODE dynamics and stochastic disturbances. This has consequences for both the estimation of functional coefficients and presenting a variance decomposition from the resulting model. We will further illustrate the application of these methods to provide linear approximations to non-linear ODE's and some consequences for approaches to registration.

REFERENCES

- [1] Frédéric Ferraty and Philippe Vieu. *Nonparametric Functional Data Analysis: Theory and Practice*. Springer, New York, 2006.
- [2] Giles Hooker and Han Lin Shang. Selecting the derivative of a functional covariate in scalar-on-function regression. *Statistics and Computing*, 32(3):35, 2022.
- [3] J Steve Marron and Andrés M Alonso. Overview of object oriented data analysis. *Biometrical Journal*, 56(5):732–753, 2014.
- [4] James O. Ramsay. Principal Differential Analysis: Data Reduction by Differential Operators. *Journal of the Royal Statistical Society. Series B (Methodological)*, 58(3):495–508, 1996. Publisher: [Royal Statistical Society, Wiley].
- [5] James O. Ramsay and B. W. Silverman. *Functional Data Analysis*. Springer Series in Statistics. Springer-Verlag, New York, 2 edition, 2005.

DEPARTMENT OF STATISTICS AND DATA SCIENCE, UNIVERSITY OF PENNSYLVANIA
Email address: ghooker@wharton.upenn.edu

ESTIMATING THE SPECTRAL DENSITY OF A FUNCTION-VALUED SPATIAL PROCESS

TAILEN HSING, RAFAIL KARTSIOUKAS, STILIAN STOEV

Classification AMS 2020: Primary 62M10; secondary 62M15, 62R10, 60G10

Keywords: functional data, spatial process, spectral density, minimax rate, central limit theorem

This paper studies the nonparametric estimation of the spectral density for a continuous-time stationary process $X = \{X(t), t \in \mathbb{R}^d\}$ taking values in some real Hilbert space. One of the novelties of our work is the consideration of functional data $X(t)$ sampled at irregular spatial locations $t_1, \dots, t_n \in \mathbb{R}^d$ as opposed to at regular grid points, e.g., $t = 1, 2, \dots$, as in functional time series. In general, spatial data are not gridded data. An excellent example is provided by the Argo dataset, in which functional data are collected over the depth of the ocean at irregular spatial locations and time points; for information of the Argo data, see Roemmich et al. (2012), Kuusela and Stein (2018), Yarger et al. (2022).

There is a sizable literature that considers the spectral-domain analysis of irregularly sampled real spatial data; some recent works include Fuentes (2007), Bandyopadhyay and Lahiri (2009), Matsuda and Yajima (2009), Subba Rao (2018), and Zhang (2024), to mention a few. Spectral inference in the realm of functional data has focused on functional time series; some examples include Panaretos and Tavakoli (2013), van Delft and Eichler (2018), Kuenzer et al. (2021), Zhu and Politis (2020), and van Delft and Dette (2024). Our paper is the first that considers spatial functional data sampled at irregular locations.

We consider in this paper a so-called lag-window estimator (cf. Brockwell and Davis, 2006; Zhu and Politis, 2020) based on estimating the covariance, which can accommodate quite general observational schemes. The performance of the estimator will be evaluated by asymptotic theory. In doing so, we assume the framework of the so-called mixed-domain asymptotics (Hall and Patil, 1994; Matsuda and Yajima, 2009; Subba Rao, 2018), which means that the sampling locations become increasingly dense and the sampling region becomes increasingly large as the number of observations increases. The rate bound of the mean squared error of our estimator is developed for a rather general mixed-domain setting. However, when data are observed on a regular grid assuming a specific covariance model, the rate bound calculations can be made rather precise, paving the way for assessing the optimality of the estimator. In particular, we establish the rate minimaxity of our estimator based on gridded data if the decay of the covariance function is dominated by a power law. The asymptotic normality of the spectral density estimator is also established under general conditions for Gaussian Hilbert-space valued processes. Finally, with a view towards practical

Date: May 1, 2025.

applications the asymptotic results are specialized to the case of discretely-sampled functional data in a reproducing kernel Hilbert space.

REFERENCES

- Bandyopadhyay, S. and Lahiri, S. N. (2009). Asymptotic properties of discrete fourier transforms for spatial data. *Sankhyā: The Indian Journal of Statistics, Series A (2008-)*, pages 221–259.
- Fuentes, M. (2007). Approximate likelihood for large irregularly spaced spatial data. *Journal of the American Statistical Association*, 102(477):321–331.
- Hall, P. and Patil, P. (1994). Properties of nonparametric estimators of autocovariance for stationary random fields. *Probability Theory and Related Fields*, 99(3):399–424.
- Kuenzer, T., Hörmann, S., and Kokoszka, P. (2021). Principal component analysis of spatially indexed functions. *Journal of the American Statistical Association*, 116(535):1444–1456.
- Kuusela, M. and Stein, M. (2018). Locally stationary spatio-temporal interpolation of Argo profiling float data. *Proceedings of the Royal Society A: Mathematical, Physical and Engineering Sciences*, 474(2220):20180400.
- Matsuda, Y. and Yajima, Y. (2009). Fourier analysis of irregularly spaced data on $\mathbb{R}^d$. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)*, 71(1):191–217.
- Panaretos, V. M. and Tavakoli, S. (2013). Fourier analysis of stationary time series in function space. *The Annals of Statistics*, 41(2):568–603.
- Roemmich, D., Gould, W. J., and Gilson, J. (2012). 135 years of global ocean warming between the Challenger expedition and the Argo Programme. *Nature Climate Change*, 2(6):425–428.
- Subba Rao, S. (2018). Statistical inference for spatial statistics defined in the Fourier domain. *Ann. Statist.*, 46(2):469–499.
- van Delft, A. and Dette, H. (2024). A general framework to quantify deviations from structural assumptions in the analysis of nonstationary function-valued processes. *The Annals of Statistics*, 52(2):550–579.
- van Delft, A. and Eichler, M. (2018). Locally stationary functional time series. *Electronic Journal of Statistics*, 12(1):107 – 170.
- Yarger, D., Stoev, S., and Hsing, T. (2022). A functional-data approach to the Argo data. *The Annals of Applied Statistics*, 16(1):216–246.
- Zhang, S. (2024). Statistical analysis of irregularly spaced spatial data in frequency domain. *Journal of Time Series Analysis*, 45(5):714–738.
- Zhu, T. and Politis, D. N. (2020). Higher-order accurate spectral density estimation of functional time series. *Journal of Time Series Analysis*, 41(1):3–20.

DEPARTMENT OF STATISTICS, UNIVERSITY OF MICHIGAN, ANN ARBOR, MI, USA
Email address: thsing@umich.edu

CONTINUOUS NORMALIZING FLOWS FOR LEARNING PROBABILITY DISTRIBUTIONS AND SHORT COURSE ON STATISTICAL DEEP LEARNING

JIAN HUANG

Classification AMS 2020: 62G05, 68T07

Keywords: conditional distribution, Deep neural networks, Nonasymptotic error bounds, Non-parametric estimation, Drift and score functions

The IMS program, held at the National University of Singapore from August 20 to September 10, 2024, focused on the forefront of functional data analysis, particularly in the context of artificial intelligence. This event brought together experts from diverse fields to explore the challenges and opportunities arising from the integration of AI into functional data analysis. As part of the program, I delivered a talk and conducted a short course, both of which focused on innovative approaches in generative learning and deep learning models.

In my talk, titled “Conditional Stochastic Interpolation for Generative Learning,” I introduced a novel method aimed at learning conditional distributions, known as Conditional Stochastic Interpolation (CSI). This method is a process-based generative model for learning conditional distributions with complex and high-dimensional data.

- *Methodology:* The CSI method is grounded in the estimation of probability flow equations or stochastic differential equations. These equations are instrumental in transporting a reference distribution to a target conditional distribution. The core of this approach lies in learning the conditional drift and score functions, which are essential for constructing deterministic processes. These processes are governed by ordinary differential equations or diffusion processes, facilitating conditional sampling.
- *Adaptive Diffusion Term:* A significant innovation in our approach is the incorporation of an adaptive diffusion term. This element is crucial for addressing instability issues that often arise in diffusion processes, thereby enhancing the robustness and reliability of the method.
- *Nonparametric Regression Approach:* We derived explicit expressions for the conditional drift and score functions in terms of conditional expectations. This derivation naturally leads to a nonparametric regression approach for estimating these functions, offering a flexible and powerful tool for practitioners.
- *Theoretical Guarantees:* To ensure the reliability of the CSI method, we established nonasymptotic error bounds for learning the target conditional distribution. These bounds provide theoretical guarantees for the method’s performance, making it a robust choice for various applications.
- *Application in Image Generation:* The practical utility of CSI was demonstrated through its application in image generation using a benchmark image dataset. This example highlighted the method’s effectiveness in generating realistic images, showcasing its potential in real-world applications.

The short course I conducted was titled “Deep Learning and Generative Models: with Applications in Statistical Analysis.” This course was designed to provide participants with a comprehensive understanding of how deep learning and generative models can be leveraged for statistical analysis.

- (1) *Deep Neural Networks for High-Dimensional Function Approximation*: This section of the course focused on the utilization of deep neural networks to approximate high-dimensional functions within nonparametric statistical models. By laying the groundwork for understanding the power of neural networks, participants gained insights into how these models can address complex statistical modeling challenges.
- (2) *Generative Learning Approaches*: We explored various generative learning approaches for modeling probability and conditional distributions. This included an in-depth look at generative adversarial networks (GANs), diffusion models, and continuous normalizing flows. These methods have demonstrated remarkable success in generating realistic data samples and understanding underlying data distributions, making them invaluable tools for statisticians and data scientists.
- (3) *Leveraging Large Models for Statistical Analysis*: The course also examined how large models can be harnessed to support statistical analysis. We covered techniques for utilizing the extensive capabilities of large models to improve accuracy and efficiency in statistical learning tasks. This section emphasized the transformative potential of large models in enhancing statistical methodologies.
- (4) *Real-World Applications*: Throughout the course, a variety of datasets were used to illustrate the broad applicability of these methods in real-world scenarios. Examples included conditional sample generation, protein sequence analysis, and image generation, demonstrating the versatility and effectiveness of deep learning and generative models in diverse fields.

DEPARTMENTS OF DATA SCIENCE AND AI, AND APPLIED MATHEMATICS, THE HONG KONG POLYTECHNIC UNIVERSITY

Email address: j.huang@polyu.edu.hk

NEURAL TANGENT KERNEL IN IMPLIED VOLATILITY FORECASTING: A NONLINEAR FUNCTIONAL AUTOREGRESSION APPROACH

HANNAH L. H. LAI

Classification AMS 2020: 46T99, 68T01

Keywords: Nonlinear Functional Autoregression; Neural Tangent Kernel; Implied Volatility Forecasting.

We denote by $\mathcal{H} = L^2(\mathcal{I})$ the Hilbert space consisting of all square-integrable surfaces defined on a compact set $\mathcal{I} \subset \mathbb{R}^q$ and equipped with the inner product $\langle f, g \rangle_{\mathcal{H}} = \int_{\mathcal{I}} f(u)g(u) du$, for any $f, g \in L^2(\mathcal{I})$. Define the squared L^2 norm of a function by $\|f\|_{\mathcal{H}}^2 = \langle f, f \rangle_{\mathcal{H}}$.

Let $\{Y_i\}_{i=1}^n$ be a series of n random surfaces that take values on $\mathcal{H}_Y = L^2(\mathcal{I}_Y)$. Associated with each Y_i , there is a regressor surface $X_i \in \mathcal{H}_X = L^2(\mathcal{I}_X)$. We consider functions with finite second moment, i.e., $\mathbb{E}[\|Y_i\|_{\mathcal{H}_Y}^2] < \infty$ and $\mathbb{E}[\|X_i\|_{\mathcal{H}_X}^2] < \infty$. For simplicity, we assume that Y_i and X_i are centered functions, i.e., $\mu_X(v) = \mathbb{E}[X_i(v)] = 0, \forall v \in \mathcal{I}_X$ and $\mu_Y(u) = \mathbb{E}[Y_i(u)] = 0, \forall u \in \mathcal{I}_Y$. Let P_X and P_Y denote the distributions of X and Y , and $P_{Y|X} : \mathcal{H}_X \times \mathcal{H}_Y \rightarrow \mathbb{R}$ the conditional distribution of Y given X . If $L_2(P_X)$ represents the class of all measurable functions of X with $\mathbb{E}[f^2(X)] < \infty$ under P_X , then $L_2(P_Y)$ is similarly defined for Y . Our goal is to capture the potential nonlinear dependence between Y_i and X_i through a function $g : \mathcal{H}_X \rightarrow \mathcal{H}_Y$

$$(0.1) \quad Y_i = g(X_i) + \epsilon_i,$$

where ϵ_i is a noise function with $\mathbb{E}[\epsilon_i(u)] = 0, \forall u \in \mathcal{I}_Y$ and $\mathbb{E}[\|\epsilon_i\|_{\mathcal{H}_Y}^2] < \infty$. In our study, X_i is a vector of lagged surfaces $Y_{i-1}, Y_{i-2}, \dots$ or their linear combination. Hence, the model (0.1) is a nonlinear functional autoregression model (NFAR).

We project Y_i onto a set of orthonormal basis functions $\varphi = (\varphi_1, \varphi_2, \dots)^T$ with $\varphi_j \in \mathcal{H}_Y$

$$(0.2) \quad Y_i = \sum_{j=1}^{\infty} y_{ij} \varphi_j, \quad \text{with } y_{ij} = \langle Y_i, \varphi_j \rangle_{\mathcal{H}_Y},$$

with $\mathbf{y}_i = (y_{i1}, y_{i2}, \dots)^T \in \mathcal{H}_y \subseteq \mathbb{R}^{\infty}$ the projection coefficients of Y_i onto the basis functions φ , satisfying $\mathbb{E}[y_{ij}y_{rv}] = 0$ for $j \neq v, j, v \in \mathbb{N}_+$ and any $i, r \in \{1, \dots, n\}$. Similarly, we project X_i onto a sequence of orthogonal basis functions $\psi = (\psi_1, \psi_2, \dots)^T$ with $\psi_j \in \mathcal{H}_X$

$$(0.3) \quad X_i = \sum_{j=1}^{\infty} x_{ij} \psi_j, \quad \text{with } x_{ij} = \langle X_i, \psi_j \rangle_{\mathcal{H}_X},$$

with $\mathbf{x}_i = (x_{i1}, x_{i2}, \dots)^T \in \mathcal{H}_x \subseteq \mathbb{R}^{\infty}$ the projection coefficients of X_i onto the basis functions ψ , satisfying $\mathbb{E}[x_{ij}x_{rv}] = 0$ for $j \neq v, j, v \in \mathbb{N}_+$ and any $i, r \in \{1, \dots, n\}$. Transitioning from functions to vectors, we define $f : \mathcal{H}_x \rightarrow \mathcal{H}_y$

$$(0.4) \quad \mathbf{y}_i = f(\mathbf{x}_i) + \epsilon_i,$$

where ϵ_i is a noise vector with $\mathbb{E}[\epsilon_{ij}] = 0$ and $\mathbb{E}[\|\epsilon_i\|^2] < \infty$. Although vectors offer a more compact representation of functions, they still exist within an infinite-dimensional framework unless additional restrictions are assumed to hold. This inherent complexity makes the empirical estimation of Equation (0.4) challenging when working with finite sample sizes. To address this issue, we employ classical sieve methods leading to finite-dimensional vector spaces.¹

To elucidate the nonlinear relation between X_i and Y_i in Equation (0.1), we introduce another Hilbert space of functions generated by a positive-definite kernel $K : \mathcal{H}_X \times \mathcal{H}_X \rightarrow \mathbb{R}$ defined on the inner product of $\mathcal{H}_X$ through a function $\rho : \mathbb{R}^3 \rightarrow \mathbb{R}^+$, such that

$$(0.5) \quad K(X_i, X_j) = \rho(\langle X_i, X_i \rangle_{\mathcal{H}_X}, \langle X_i, X_j \rangle_{\mathcal{H}_X}, \langle X_j, X_j \rangle_{\mathcal{H}_X}),$$

for any $X_i, X_j \in \mathcal{H}_X$. The function-on-function regression problem in Equation (0.1) can be reformulated as a functional kernel regression, in which the task is to find $B_0 \in \mathcal{B}(\mathcal{H}_Y, \mathfrak{M}_X)$ such that

$$(0.6) \quad B_0 = \arg \min_{B \in \mathcal{B}(\mathcal{H}_Y, \mathfrak{M}_X)} \mathbb{E}[\|Y_i - B^* K(., X_i)\|_{\mathcal{H}_Y}^2].$$

The solution for the kernel functional regression can be found in [2]. We define a new kernel $k : \mathcal{H}_x \times \mathcal{H}_x \rightarrow \mathbb{R}$ such that for any $\mathbf{x}_i, \mathbf{x}_j \in \mathcal{H}_x$

$$(0.7) \quad k(\mathbf{x}_i, \mathbf{x}_j) = \rho(\langle \mathbf{x}_i, \mathbf{x}_i \rangle, \langle \mathbf{x}_i, \mathbf{x}_j \rangle, \langle \mathbf{x}_j, \mathbf{x}_j \rangle).$$

Lemma 0.1 (Isomorphism between Reproducing Kernel Hilbert Spaces). *Under Equations (0.3) and (0.7), it holds that*

$$(0.8) \quad \begin{aligned} k(\mathbf{x}_i, \mathbf{x}_j) &= \langle k(., \mathbf{x}_i), k(., \mathbf{x}_j) \rangle \\ &= \langle K(., X_i), K(., X_j) \rangle_{\mathfrak{M}_X} = K(X_i, X_j). \end{aligned}$$

Then the RKHS $\mathfrak{M}_X$ nested on $\mathcal{H}_X$ is isometrically isomorphic to the RKHS $\mathfrak{M}_x$ nested on $\mathcal{H}_x$.

Theorem 0.2 (Vector-to-vector regression). *Given the decomposition of Y_i in Equation (0.2) and X_i in Equations (0.3), under some technical Assumptions and Lemma 0.1, for a positive definite kernel k defined by Equation (0.7), if there is a covariance matrix Σ_{xx} of $k(., \mathbf{x})$ that is diagonal, then the function-to-function regression model in Equation (0.6) may be represented equivalently by*

$$(0.9) \quad \beta_0 = \arg \min_{\beta \in \mathcal{B}(\mathcal{H}_Y, \mathfrak{M}_x)} \mathbb{E}[\|\mathbf{y}_i - \beta^* k(., \mathbf{x}_i)\|^2],$$

with solution $\beta_0 = \Sigma_{xx}^\dagger \Sigma_{xy}$. This leads to

$$(0.10) \quad \begin{aligned} \mathbb{E}[\mathbf{y}_i | \mathbf{x}_i] &= \beta_0^* k(., \mathbf{x}_i) \\ &= \Sigma_{yx} \Sigma_{xx}^\dagger k(., \mathbf{x}_i) \\ &= \mathbb{E}[\{(\Sigma_{xx}^\dagger k(., \mathbf{x}_i))(\mathbf{x})\} \mathbf{y}]. \end{aligned}$$

¹Sieve methods involve truncating the regression for the full set of projection coefficients while striving to minimize any loss of information.

In our work, we utilize the Neural Tangent Kernel (NTK) of [1], a flexible kernel class that uses neural networks to capture complex nonlinear dependencies in data effectively. The NTK describes how neural networks behave under first-order gradient descent training and is calculated as the inner product of the network's weight gradients. Our empirical analysis includes over 6 million European calls and put options from the S&P 500 Index, covering January 2009 to December 2021. The results confirm the superior forecasting accuracy of the fNTK across different time horizons. When applied to short delta-neutral straddle trading, the fNTK achieves a Sharpe ratio ranging from 1.30 to 1.83 on a weekly to monthly basis, translating to 90% to 675% relative improvement in portfolio returns compared to forecasts based on functional Random Walk model.

REFERENCES

- [1] Jacot, A., F. Gabriel, and C. Hongler . Neural Tangent Kernel: Convergence and Generalization in Neural Networks. *Advances in Neural Information Processing Systems*, 31, 2018.
- [2] Sang, P. and B. L. Nonlinear Function-on-Function Regression by RKHS. *arXiv preprint arXiv:2207.08211*, 2022.

Email address: hlhlai@u.nus.edu

ROBUST MODEL AVERAGING PREDICTION

JIALIANG LI

Classification AMS 2020:

Keywords: Longitudinal data analysis, Microbiome data analysis, Model averaging, Rank regression, Robust estimation, Sure screening

Model averaging is an attractive ensemble technique to construct fast and accurate prediction. Despite of having been widely practiced in cross-sectional data analysis, its application to longitudinal data is rather limited so far. We consider model averaging for longitudinal response when the number of covariates is ultrahigh. To this end, we propose a novel two-stage procedure in which variable screening is first conducted and then followed by model averaging. In both stages, a robust rank-based estimation function is introduced to cope with potential outliers and heavy-tailed error distributions, while the longitudinal correlation is modeled by a modified Cholesky decomposition method and properly incorporated to achieve efficiency. Asymptotic properties of our proposed methods are rigorously established, including screening consistency and convergence of the model averaging estimates. Extensive simulation studies demonstrate that our method outperforms existing competitors, resulting in significant improvements in screening and prediction performance. Finally, we apply our proposed framework to analyze a human microbiome dataset, showing the capability of our procedure in resolving robust prediction using massive metabolites.

DEPARTMENT OF STATISTICS & DATA SCIENCE, NATIONAL UNIVERSITY OF SINGAPORE, SINGAPORE
Email address: jialiang@nus.edu.sg

FUNCTIONAL PRINCIPAL COMPONENT ANALYSIS FOR DISTRIBUTION-VALUED PROCESSES

HANG ZHOU AND HANS-GEORG MÜLLER

Keywords: Distributional Data Analysis, Functional Data Analysis, Longitudinal Data Analysis, Sparse Designs, Stochastic Process, Wasserstein Metric

Functional data are samples of realizations of square integrable scalar or vector-valued functions that have been extensively studied (Hsing & Eubank 2015, Wang et al. 2016). The restriction to the realm of Euclidean space-valued functions that also encompasses Hilbert-space valued functional data, i.e., function-valued stochastic processes (Chen et al. 2017), is an essential feature of functional data, but proves too restrictive as new complex non-Euclidean data types are emerging. A previous very general model for the case of a metric-space valued process for which one observes a sample of realizations (Dubey & Müller 2020) includes distribution-valued processes as a special case. The general framework developed in this previous approach utilizes a notion of metric covariance and includes a certain kind of functional principal component analysis for general metric space-valued processes by using Fréchet integrals (Petersen & Müller 2016) and is limited to the case of fully observed metric space-valued functional data, where it is assumed that $X_i(t)$ is known for all t in the time domain and cannot be extended to the case of sparsely sampled processes. We present models and analysis tools for a specific yet important class of random object-valued stochastic processes: those where the time-indexed objects are univariate distributions. Distribution-valued stochastic processes are encountered in various complex applications. We start with an i.i.d. sample of realizations of such processes. The statistical modeling of distribution-valued processes is an essential yet still missing tool for the emerging field of distributional data analysis (Petersen et al. 2022), while various modeling approaches for distributional regression and distributional time series have been studied recently (Kokoszka et al. 2019, Ghodrati & Panaretos 2022, Chen et al. 2023, Zhu & Müller 2023).

We aim for intrinsic modeling of distributions rather than extrinsic approaches. An issue that is of additional practical relevance and theoretical interest is that available observations typically are not available continuously in time but only at discrete time points. These considerations motivate a comprehensive intrinsic model for distribution-valued processes where the processes may be fully or only partially observed. Throughout we work with the 2-Wasserstein metric $d_{W,2}$ and optimal transports, which move distributions along geodesics. The challenge of intrinsic modeling is that the Wasserstein space of distributions does not have a linear or vector space structure. This challenge can be addressed by making use of rudimentary algebraic operations on the space of optimal transports (Zhu & Müller 2023). From the outset we aim to deal with centered processes. Since no subtraction exists in the Wasserstein space, the centering of distribution-valued processes is achieved by substituting transport processes for distributional processes: For each time argument the distributions that constitute the values of a distributional process at a fixed time t are replaced by

transports from the barycenter (Fréchet mean) of the process at t to the distribution that corresponds to the value of the process at time t . These transports are well defined if one adopts the Wasserstein metric. Their Fréchet mean is the identity transport, i.e., these transports are centered.

For our study of stochastic transport processes we introduce representations

$$T(t) = g(Z(t)) \odot T_0,$$

where $Z(t)$ is a $\mathbb{R}$ -valued random process, g is a bijective function that maps $\mathbb{R}$ to $(-1, 1)$ and T_0 is a single random transport that is a summary characteristic for each realization of the transport process. Here $\odot$ is a multiplication operation by which a transport is multiplied with a scalar (Zhu & Müller 2023). By construction, $g(Z(t)) \odot T_0$ lies on the extended geodesic that passes through T_0 . We develop a predictor for each individual $T_i(t)$ based on observations obtained at discrete time points and establish asymptotic convergence rates for the components of the model for both densely and sparsely sampled distributional processes. These are novel even for classical real-valued functional data.

Let $\mathcal{W}$ be the set of finite second moment probability measures on the closed interval $\mathcal{S} \subset \mathbb{R}$,

$$\mathcal{W} = \left\{ \mu \in \mathcal{P}(\mathcal{S}) : \int_{\mathcal{S}} |x|^2 d\mu(x) < \infty \right\}, \quad (1)$$

where $\mathcal{P}(\mathcal{S})$ is the set of all probability measures on $\mathcal{S}$. The p -Wasserstein distance $d_{W,p}(\cdot, \cdot)$ between two measures $\mu, \nu \in \mathcal{W}$ is

$$d_{W,p}(\mu, \nu) := \inf \left\{ \left(\int_{\mathcal{S}^2} |x_1 - x_2|^p d\Gamma(x_1, x_2) \right)^{1/p} : \Gamma \in \Gamma(\mu, \nu) \right\} \quad \text{for } p > 0, \quad (2)$$

where $\Gamma(\mu, \nu)$ is the set of joint probability measures on $\mathcal{S}^2$ with μ and ν as marginal measures. The Wasserstein space $(\mathcal{W}, d_{W,p})$ is a separable and complete metric space (Ambrosio et al. 2008, Villani et al. 2009). Here we assume $\mathcal{S} = [0, 1]$ without loss of generality to simplify the notation. Given two probability measures $\mu, \nu \in \mathcal{W}$, the optimal transport from μ to ν is the map $T : \mathcal{S} \rightarrow \mathcal{S}$ that minimizes the transport cost,

$$\arg \inf_{T \in \mathcal{T}} \left\{ \left(\int_{\mathcal{S}} |T(u) - u|^p d\mu(u) \right)^{1/p}, \text{ such that } T\#\mu = \nu \right\}, \quad (3)$$

where $\mathcal{T} = \{T : \mathcal{S} \mapsto \mathcal{S} | T(0) = 0, T(1) = 1, T \text{ is non-decreasing}\}$ is the transport space and $T\#\mu$ is the push-forward measure of μ , defined as $(T\#\mu)(A) = \mu\{x \in \mathcal{S} | T(x) \in A\}$ for all A in the Borel algebra of $\mathcal{S}$. This optimization problem, also known as the Monge problem, is a relaxation of the Kantorovich problem (2). If μ is absolutely continuous with respect to the Lebesgue measure, then problems (2) and (3) are equivalent and have a unique solution $T(u) = F_\nu^{-1} \circ F_\mu(u)$ for $p = 2$, where F_μ and F_ν^{-1} are the cumulative distribution and quantile functions of μ and ν , respectively (Gangbo & McCann 1996).

For a distribution-valued process $X(t)$ with random distributions on domain $\mathcal{S}$ where $t \in \mathcal{D}$ for a closed interval in $\mathbb{R}$, the cross-sectional Fréchet mean of $X(t)$ at each t is

$$\mu_{\oplus,2}(t) = \operatorname{argmin}_{\omega \in \mathcal{W}} \mathbb{E} d_{W,2}^2(X(t), \omega).$$

We then define the (optimal) transport process $T(\cdot)$, where $T(t)$ represents the optimal transport from $\mu_{\oplus,2}(t)$ to $X(t)$, $\mu_{\oplus,2}(t)$ serves as the mean, and the transport $T(t)$ from $\mu_{\oplus,2}(t)$ to $X(t)$ quantifies the difference between $X(t)$ and $\mu_{\oplus,2}(t)$ for each $t \in \mathcal{D}$ under the Wasserstein metric. It is thus advantageous to use the transport space $\mathcal{T}$.

A scalar multiplication operation in the transport space (Zhu & Müller 2023),

$$\alpha \odot T(u) := \begin{cases} u + \alpha\{T(u) - u\}, & 0 < \alpha \leq 1 \\ u, & \alpha = 0 \\ u + \alpha\{u - T^{-1}(u)\}, & -1 \leq \alpha < 0 \end{cases}$$

induces a geodesic on $\mathcal{T}$ from $\operatorname{Unif}_{\mathcal{S}}$ to T , denoted by $u \odot T$ for all $u \in [-1, 1]$. We introduce a binary relation $\sim$ on $\mathcal{T}$, defined as $T_1 \sim T_2$ if and only if there exists $a \in [0, 1]$ such that $T_1 = a \odot T_2$ or $T_2 = a \odot T_1$ and demonstrate that $\sim$ is an equivalence relation on $\mathcal{T}$.

In analogy to the decomposition of Euclidean-valued functional data into a mean function and a stochastic part, we assume that the centered transport processes $T(t)$ can be decomposed into a scalar random function $U(t)$ that serves as a scalar multiplier in the transport space and a characteristic overall transport T_0 ,

$$T(t) = U(t) \odot T_0, \text{ for all } t \in \mathcal{D}, \quad (4)$$

where T_0 is a random element in $\mathcal{T}$ associated with each realization of the transport process. The scalar multiplier function is a stochastic process that takes values in $(-1, 1)$ and is derived from an unconstrained process Z through a transformation g as follows,

$$U(t) = g(Z(t)), \quad Z(t) \in \mathbb{R}, \quad \mathbb{E}[Z(t)] = 0, \quad g : \mathbb{R} \mapsto (-1, 1), \quad g \text{ is bijective, for all } t \in \mathcal{D}. \quad (5)$$

The mean zero stochastic process $Z(t)$ in conjunction with the bijective map $g : \mathbb{R} \mapsto (-1, 1)$ further characterizes the transport process T , where $T(t)$ resides in $\{T : T \in [T_0]_{\sim}\} \cup \{T : T \in [T_0^{-1}]_{\sim}\}$, which includes the geodesic from T_0^{-1} to T_0 .

For some situations it is appropriate and advantageous to further assume that the process Z is a Gaussian process, a property that can be harnessed to obtain methods for the important case where the distribution-valued trajectories are only observed on a discrete grid of time points that might be sparse. that the stochastic transport process (4) is well-defined.

Further details can be found in the preprint: Zhou, H. and Müller, H.G., 2023. Optimal transport representations and functional principal components for distribution-valued processes. arXiv preprint arXiv:2310.20088.

References

- Ambrosio, L., Gigli, N. & Savaré, G. (2008), *Gradient Flows: in Metric Spaces and in the Space of Probability Measures*, Springer Science & Business Media.
- Chen, K., Delicado, P. & Müller, H.-G. (2017), ‘Modeling function-valued stochastic processes, with applications to fertility dynamics’, *J. R. Stat. Soc. Ser. B Stat. Methodol.* **79**, 177–196.
- Chen, Y., Lin, Z. & Müller, H.-G. (2023), ‘Wasserstein regression’, *J. Amer. Statist. Assoc.* **118**(542), 869–882.
- Dubey, P. & Müller, H.-G. (2020), ‘Functional models for time-varying random objects’, *J. R. Stat. Soc. Ser. B Stat. Methodol.* **82**(2), 275–327.
- Gangbo, W. & McCann, R. J. (1996), ‘The geometry of optimal transportation’, *Acta Math.* **177**, 113–161.
- Ghodrati, L. & Panaretos, V. M. (2022), ‘Distribution-on-distribution regression via optimal transport maps’, *Biometrika* **109**(4), 957–974.
- Hsing, T. & Eubank, R. (2015), *Theoretical Foundations of Functional Data Analysis, with an Introduction to Linear Operators*, John Wiley & Sons.
- Kokoszka, P., Miao, H., Petersen, A. & Shang, H. L. (2019), ‘Forecasting of density functions with an application to cross-sectional and intraday returns’, *Int. J. Forecast.* **35**(4), 1304–1317.
- Petersen, A. & Müller, H.-G. (2016), ‘Fréchet integration and adaptive metric selection for interpretable covariances of multivariate functional data’, *Biometrika* **103**(1), 103–120.
- Petersen, A., Zhang, C. & Kokoszka, P. (2022), ‘Modeling probability density functions as data objects’, *Econom. Stat.* **21**, 159–178.
- Villani, C. et al. (2009), *Optimal Transport: Old and New*, Springer.
- Wang, J.-L., Chiou, J.-M. & Müller, H.-G. (2016), ‘Functional data analysis’, *Annu. Rev. Stat. Appl.* **3**, 257–295.
- Zhu, C. & Müller, H.-G. (2023), ‘Autoregressive optimal transport models’, *J. R. Stat. Soc. Ser. B Stat. Methodol.* **85**(3), 1012–1033.

HIGH-DIMENSIONAL HILBERT-SCHMIDT LINEAR REGRESSION FOR HILBERT MANIFOLD VARIABLES

BYEONG UK PARK

Classification AMS 2020: 62R30, 62J07.

Keywords: non-Euclidean data, high-dimensional regression, Hilbert-Schmidt operators, spectral decomposition, penalization.

We present a unified framework for high-dimensional linear regression with non-Euclidean data where both the response variable and covariates take values in general Riemannian Hilbert manifolds. In our modeling, the response and covariates are allowed to originate from distinct spaces and are interconnected by Hilbert-Schmidt operators. The methodology is developed under a general penalization scheme incorporating various non-convex penalty functions, thereby accommodating scenarios where the number of covariates grows exponentially with the sample size. Leveraging modern statistical theory for data residing on Hilbert manifolds, we establish the oracle property and derive error bounds for the proposed estimators. The practical validity of the proposed method is demonstrated via numerical simulation and real data applications.

Specifically, let Y and X_j for $1 \leq j \leq p$ be random variables taking values in Riemannian Hilbert manifolds $\mathcal{M}_Y$ and $\mathcal{M}_j$ respectively. Let $D_j := \dim(\mathcal{M}_j)$ be the dimensions of $\mathcal{M}_j$. We allow the case where $\dim(\mathcal{M}_Y) = \infty$. Let Log_y^Y and $\text{Log}_{x_j}^j$ be Riemannian logarithmic maps at y and x_j in $\mathcal{M}_Y$ and $\mathcal{M}_j$, respectively. Also, let μ_Y and μ_j , respectively, denote the Fréchet means of Y and X_j . We consider the following Hilbert-Schmidt linear model:

$$(0.1) \quad \text{Log}_{\mu_Y}^Y Y = \sum_{j=1}^p \mathfrak{B}_j(\text{Log}_{\mu_j}^j X_j) + \varepsilon,$$

where $\mathfrak{B}_j$ are Hilbert-Schmidt operators and ε is a random error. We assume a high-dimensional setting where p diverges as the sample size n grows. In this setting, we impose a sparsity condition on the Hilbert-Schmidt operators $\mathfrak{B}_j$, meaning that the number of nonzero operators is relatively small. Our primary goal is to estimate the operators $\mathfrak{B}_j$ and recover the index set $\mathcal{S} := \{1 \leq j \leq p : \mathfrak{B}_j \neq 0\}$.

The estimation procedure is based on the spectral decomposition for X_j . Let $\hat{\mu}_Y$ and $\hat{\mu}_j$ be the empirical Fréchet means corresponding to μ_Y and μ_j , respectively. Then, the empirical covariance operators are given by $\hat{\mathcal{C}}_j := n^{-1} \sum_{i=1}^n \text{Log}_{\hat{\mu}_j}^j X_{ij} \otimes \text{Log}_{\hat{\mu}_j}^j X_{ij}$. Each $\hat{\mathcal{C}}_j$ admits a spectral decomposition $\hat{\mathcal{C}}_j = \sum_{k=1}^{D_j} \hat{\omega}_{jk}(\hat{e}_{jk} \otimes \hat{e}_{jk})$ with eigenvalues $\hat{\omega}_{jk}$ and the corresponding orthonormal basis $\{\hat{e}_{jk} : 1 \leq k \leq D_j\}$. From the spectral decomposition we get the estimated k th scores $\hat{\xi}_{i,jk}$ of $\text{Log}_{\hat{\mu}_j}^j X_{ij}$. We introduce (truncation) parameters K_j that diverge to infinity as the sample size increases for infinite-dimensional $\mathcal{M}_j$, and

are equal to D_j for finite-dimensional $\mathcal{M}_j$. Then, under the model (0.1) it holds that

$$\text{Log}_{\hat{\mu}_Y}^Y Y_i \approx \sum_{j=1}^p \sum_{k=1}^{K_j} \hat{\xi}_{i,jk} \cdot \mathfrak{B}_j^*(\hat{e}_{jk}) + \mathcal{P}_{\mu_Y, \hat{\mu}_Y}^Y(\varepsilon_i)$$

where $\mathfrak{B}_j^* := \mathcal{P}_{\mu_Y, \hat{\mu}_Y}^Y \circ \mathfrak{B}_j \circ \mathcal{P}_{\hat{\mu}_j, \mu_j}^j$ and $\mathcal{P}_{\mu_Y, \hat{\mu}_Y}^Y(\mathcal{P}_{\hat{\mu}_j, \mu_j}^j)$ is the parallel transport that maps the tangent space $T_{\mu_Y} \mathcal{M}_Y$ ($T_{\hat{\mu}_j} \mathcal{M}_j$) of $\mathcal{M}_Y$ ($\mathcal{M}_j$) at μ_Y ($\hat{\mu}_j$) to the tangent space $T_{\hat{\mu}_Y} \mathcal{M}_Y$ ($T_{\mu_j} \mathcal{M}_j$) at $\hat{\mu}_Y$ (μ_j).

We actually estimate $\mathfrak{B}_j^*$ instead of $\mathfrak{B}_j$. Let $\beta_{jk}^* := \hat{\omega}_{jk}^{1/2} \mathfrak{B}_j^*(\hat{e}_{jk})$. We consider a general class of penalty functions ρ_λ , which encompasses the LASSO, SCAD and MCP penalty functions. With $\lambda_j := \sqrt{K_j} \cdot \lambda$ for a universal penalty parameter $\lambda > 0$, we formulate a penalized objective function $\mathcal{L}_n$ defined by

$$\mathcal{L}_n(\beta) := \frac{1}{2n} \sum_{i=1}^n \left\| \text{Log}_{\hat{\mu}_Y}^Y Y_i - \sum_{j=1}^p \sum_{k=1}^{K_j} \hat{\xi}_{i,jk} \hat{\omega}_{jk}^{-1/2} \cdot \beta_{jk} \right\|^2 + \sum_{j=1}^p \rho_{\lambda_j}(\|\beta_j\|),$$

where $\beta_j = (\beta_{j1}, \dots, \beta_{jK_j})^\top \in (T_{\hat{\mu}_Y} \mathcal{M}_Y)^{K_j}$ and $\beta = (\beta_1^\top, \dots, \beta_p^\top)^\top$. We solve the following constrained minimization problem:

$$(0.2) \quad \hat{\beta}^* := \arg \min \left\{ \mathcal{L}_n(\beta) : \beta \in (T_{\hat{\mu}_Y} \mathcal{M}_Y)^{K_+} \text{ with } \sum_{j=1}^p \sqrt{K_j} \|\beta_j\| \leq R \right\}$$

for some regularization parameter $R \geq 0$. The Hilbert-Schmidt operators $\mathfrak{B}_j^*$ are then estimated by

$$\hat{\mathfrak{B}}_j^* := \sum_{k=1}^{K_j} \hat{\omega}_{jk}^{-1/2} \cdot (\hat{e}_{jk} \otimes \hat{\beta}_{jk}^*).$$

We study the statistical properties of the estimators $\hat{\beta}^*$ and the corresponding $\hat{\mathfrak{B}}_j^*$. We first derive the rates of convergence of the eigenvalues $\hat{\omega}_{jk}$ and the corresponding orthonormal bases $\hat{e}_{jk}$ to their population counterparts that are uniform over the diverging number of covariates. For this, we elicit some concentration inequalities for the empirical Fréchet means $\hat{\mu}_j$ using empirical process theory. Built on these results, we derive the estimation error bounds of a stationary solution $\hat{\beta}^*$ of the constrained minimization problem (0.2) in various convergence modes. We also show that $\hat{\beta}^*$ exhibits the oracle property with selection consistency. From these results for $\hat{\beta}^*$ we establish the error bounds and the oracle property of the estimated Hilbert-Schmidt operators $\hat{\mathfrak{B}}_j^*$.

REFERENCES

- [1] Changwon Choi, and Byeong U. Park (2024). High-dimensional Hilbert-Schmidt linear regression for Hilbert manifold variables. *Manuscript*.

DEPARTMENT OF STATISTICS, SEOUL NATIONAL UNIVERSITY
Email address: bupark@snu.ac.kr

LARGE-SCALE MULTIPLE TESTING OF CROSS-COVARIANCE FUNCTIONS WITH APPLICATIONS TO FUNCTIONAL NETWORK MODELS

XINGHAO QIAO

Classification AMS 2020: 62A01, 62H15.

Keywords: Discretely observed functional data, False discovery control, Functional covariance model, Functional graphical model, High-dimensional functional data

Recent advances in information technology have led to the growing prevalence of multivariate or even high-dimensional functional datasets across various applications. Examples include time-course gene expression data in genomics, and various types of brain imaging data in neuroscience, such as EEG, MEG and fMRI, among others. The brain signals are typically recorded in the form of a location-by-time matrix for each individual subject, where the rows and columns correspond to a set of brain regions and a number of observational time points spanning over minutes, respectively. To capture the non-stationary and dynamic patterns, recent proposals involve modeling brain signals as multivariate random functions, treating the time course data of each region as a random function. Such high-dimensional functional data can be represented as $\mathbf{X}_i(\cdot) = \{X_{i1}(\cdot), \dots, X_{ip}(\cdot)\}^T$ defined on a compact interval $\mathcal{U}$ with marginal- and cross-covariance functions, which together form the covariance function matrix

$$\Sigma(\cdot, \cdot) = \{\Sigma_{jk}(\cdot, \cdot)\}_{p \times p}, \quad \Sigma_{jk}(u, v) = \text{Cov}\{X_{ij}(u), X_{ik}(v)\} \text{ for } (u, v) \in \mathcal{U}^2.$$

We observe $\mathbf{X}_i(\cdot)$ for $i = 1, \dots, n$, where the dimension p is large relative to, and maybe greater than the number of subjects n .

Our motivation lies in the brain connectivity analysis based on brain imaging data. The literature primarily focuses on two types of multivariate functional data analysis approaches for estimating brain connectivity networks, consisting of p nodes. The first method [1] considers a functional covariance model depicting the marginal correlation information in $\mathbf{X}_i(\cdot)$. This approach aims to identify the functional sparsity structure in $\Sigma(\cdot, \cdot)$, i.e., discovering edges (j, k) 's such that $\Sigma_{jk}(u, v) \neq 0$ for some $(u, v) \in \mathcal{U}^2$. By applying adaptive functional thresholding, they achieve estimation and support recovery consistencies. The second method frames the network estimation as a functional graphical modelling problem [2]. This model characterizes the conditional dependence structure of p Gaussian random functions, i.e., nodes j and k are connected by an edge if and only if $X_{ij}(\cdot)$ and $X_{ik}(\cdot)$ are dependent, conditional on the remaining $p - 2$ functions. This line of work has witnessed numerous recent advancements. Both types of models involve developing the regularized estimation of functional network structures that leverage the functional sparsity information. However, a precise theoretical relationship between the tuning parameter and the number of false edges is still ambiguous. Empirically, large regularization parameters result in sparse networks and may not be effective in discovering edges with small weights, small regularization parameters can produce excessive false edges, leading to high false discovery rates.

We reframe the functional covariance model estimation as a problem of simultaneously testing $p(p-1)/2$ hypotheses for cross-covariance functions:

(0.1)

$$H_{0,jk} : \Sigma_{jk}(u, v) = 0 \text{ for any } (u, v) \in \mathcal{U}^2 \text{ vs } H_{1,jk} : \Sigma_{jk}(u, v) \neq 0 \text{ for some } (u, v) \in \mathcal{U}^2,$$

where $1 \leq j < k \leq p$ and we identify a significant edge between nodes j and k if and only if $H_{0,jk}$ is rejected. In practical scenarios involving brain imaging data, where each trajectory $X_{ij}(\cdot)$ is observed at a dense set of points, nonparametric smoothing is frequently employed to obtain estimated curves $\hat{X}_{ij}(\cdot)$. This requires considering functional error-contaminated versions of $X_{ij}(\cdot)$ that satisfy:

(0.2)

$$\hat{X}_{ij}(\cdot) = X_{ij}(\cdot) + e_{ij}(\cdot), \quad i = 1, \dots, n, \quad j = 1, \dots, p.$$

Additionally, we formulate the functional graphical model estimation as a multiple testing task on the cross-covariance structure of functional regression errors, which are constructed through nodewise functional regressions. Consequently, we need to deal with estimated functional regression errors (i.e., functional residuals) instead of true ones, following a similar form as (0.2).

We aim to develop a general framework for large-scale multiple testing of cross-covariance functions from both methodological and theoretical perspectives. We begin by constructing a Hilbert–Schmidt-norm-based test statistic for each pair $j < k$, whose limiting null distribution is shown to be an infinite mixture of chi-squares. We then employ normal quantile transformations for all test statistics, upon which a multiple testing procedure is proposed to account for the multiplicity and dependence among the transformed test statistics. We establish that our procedure can control false discoveries asymptotically under both fully-observed and error-contaminated functional scenarios, and, furthermore, be applied to two concrete examples: the *functional covariance model with discrete observations* and the more important-yet-challenging *functional graphical model*. Specifically, we demonstrate how our proposed methods for both applications seamlessly integrate into the general error-contamination framework, and, with verifiable conditions, achieve theoretical guarantees on false discovery control. Empirically, we conduct simulations to showcase the uniform superiority of our proposals over competitors in terms of false discovery control and power for both fully and discretely observed functional data within both functional covariance and graphical models. We also apply our method to identify network structures using two brain imaging datasets, and observe scientifically interpretable patterns.

The main contributions of our paper are threefold. First, we make the first attempt in the literature of functional data analysis and multiple testing to develop a general procedure with theoretical guarantees for the simultaneous testing of a large collection of hypotheses for the functional sparse covariance structures. Our approach is fully functional in the sense it does not rely on dimension reduction techniques, thereby avoiding any incurred information loss. Second, we extend our method and theory to the more general error-contamination setting (0.2), encompassing the common practical scenario of discretely observed functional data as a special case. On the method front, our procedure remains valid by simply replacing each $X_{ij}(\cdot)$ with its estimated surrogate $\hat{X}_{ij}(\cdot)$. Theoretically, we demonstrate that, under an additional

condition, our proposal still ensures control over false discoveries. Such condition can be verified under mild circumstances for discretely observed functional data.

Las but not least, we develop the first effective testing procedure for high-dimensional functional graphical model inference. Specifically, we propose a novel approach to reformulate the functional graphical model estimation as a multiple testing problem for the cross-covariance structures between $p(p-1)/2$ pairs of functional regression errors formed by respectively regressing each pair $X_{ij}(\cdot)$ and $X_{ik}(\cdot)$ ($1 \leq j < k \leq p$) on the remaining $p-2$ functional variables. Employing the penalized functional regression technique to estimate functional coefficients, we obtain $p(p-1)/2$ pairs of functional residuals, which can be nicely integrated into a generalized version of our error-contamination framework. By deriving relevant convergence rates to validate the extra condition and applying our established theory, we show that such proposal combined with our multiple test procedure achieves false discovery control.

REFERENCES

- [1] Fang, Q., Guo, S. and Qiao, X. Adaptive functional thresholding for sparse covariance function estimation in high dimensions. *Journal of the American Statistical Association*, 119: 1473–1485, 2023.
- [2] Qiao, X., Guo, S. and James, G. Functional graphical models. *Journal of the American Statistical Association*, 114: 211–222, 2019.

FACULTY OF BUSINESS AND ECONOMICS, THE UNIVERSITY OF HONG KONG
Email address: xinghaoq@hku.hk

ROBUST INVERSE REGRESSION FOR MULTIVARIATE ELLIPTICAL FUNCTIONAL DATA

EFTYCHIA SOLEA

Classification AMS 2020:

Keywords: dimension reduction, elliptical distribution, functional data, sliced inverse regression, spatial sign

Functional data have received significant attention as they frequently appear in modern applications, such as functional magnetic resonance imaging (fMRI) and natural language processing. The infinite-dimensional nature of functional data makes it necessary to use dimension reduction techniques. Most existing techniques, however, rely on the covariance operator, which can be affected by heavy-tailed data and unusual observations. Therefore, in this paper, we consider a *robust functional sliced inverse regression (R-FSIR) for multivariate elliptical functional data*. For that reason, we define the elliptical distribution for a vector of random functions, extending the existing definition of [1] to the multivariate setting. We introduce a new statistical linear operator, called the *conditional spatial sign Kendall's tau covariance operator*, which can be seen as an extension of the multivariate Kendall's tau to both the conditional and functional settings, and is capable to handle heavy-tailed functional data and outliers. We show that the conditional spatial sign Kendall's tau covariance operator has the same eigenfunctions with the conditional covariance operator, and hence we can formulate the generalized eigenvalue problem based on this new operator to achieve estimation robustness. We derive the convergence rates of the proposed estimators for both completely and partially observed data. In practice, we can only observe the functions at discrete time points, and the new theoretical results support practical estimation procedure. Finally, we demonstrate the finite sample performance of our estimator using simulation examples and a real dataset based on fMRI. We observe that R-FSIR and FSIR have comparable performance for the Gaussian distribution with no outliers. However, R-FSIR outperforms FSIR for heavy-tailed data. Specifically, the efficiency of R-FSIR remains reasonably high, whereas the efficiency of FSIR decreases considerably. This is especially evident when outliers are added to the data.

REFERENCES

- [1] Boente, G., Barrera, M. S., & Tyler, D. E. (2014). A characterization of elliptical distributions and some optimality properties of principal components for functional data. *Journal of Multivariate Analysis*, 131, 254-264.
- [2] Ferre, L., and Yao, A. F. (2003). Functional sliced inverse regression analysis. *Statistics*, 37(6), 475-488.

Eftychia Solea, Eliana Christou and Jun Song Robust Inverse Regression for Multivariate Elliptical Functional Data. *Statistica Sinica*, 2024.

SCHOOL OF MATHEMATICAL SCIENCES, QUEEN MARY UNIVERSITY OF LONDON, LONDON, E1 4NS,
UNITED KINGDOM
E-mail address: e.solea@qmul.ac.uk

DEEP LEARNING FOR FUNCTIONAL AND SURVIVAL DATA

JANE-LING WANG

Classification AMS 2020: 62N02, 62G20.

Keywords: Neural network, functional data analysis, censored data, minimax theory, semi-parametric efficiency

Deep learning, aka deep neural networks, has enjoyed tremendous success in applications for all kinds of data. However, its application to functional data is limited and the theoretical foundation for why it works is still lacking. This talk explores the application of deep neural networks (DNN) to two types of data: functional data and censored survival data.

- **Functional Data:** The infinite dimensionality of functional data means standard learning algorithms can be applied only after appropriate dimension reduction, typically through basis expansions. Currently, these bases are chosen a priori without the information for the task at hand and thus may be suboptimal. We instead propose to adaptively learn these bases in an end-to-end fashion. We introduce a DNN that employs a new basis-layer whose hidden units are each basis functions themselves, implemented as a micro neural network. This architecture learns parsimonious dimension reduction to functional inputs that focuses only on information relevant to the target rather than irrelevant variation in the input function. Across numerous classification and regression tasks that involve functional data this method empirically outperforms other types of DNN.
- **Survival Data:** While DNN have demonstrated empirical success in applications for survival data, most of these methods are difficult to interpret and mathematical understanding of them is lacking. We study the partially linear Cox model, where the nonlinear component of the model is implemented using a deep neural network. The proposed approach is flexible and able to circumvent the curse of dimensionality, yet it facilitates interpretability of the effects of treatment covariates on survival. We establish asymptotic theory for maximum partial likelihood estimators and show that the nonparametric DNN estimator achieves the minimax optimal rate of convergence (up to a poly-logarithmic factor). Moreover, the corresponding parametric estimator for treatment covariate effects is $\sqrt{n}$ -consistent, asymptotically normal, and attains semiparametric efficiency. Numerical experiments provide evidence of the advantages of the proposed method.

DEPARTMENT OF STATISTICS, UNIVERSITY OF CALIFORNIA, DAVIS, CA 95616, USA
Email address: janelwang@ucdavis.edu

RECOVERY OF TIME LABELS FOR NOISY DYNAMIC DATA

WANJIE WANG

The dynamic system is frequently observed in science, with plenty of works in many directions. We observe $X_i \in R^p$ at time points t_i , $1 \leq i \leq N$. Based on X_i , we want to reconstruct the system $X(t)$. However, in some applications, the time labels t_i 's are not available. For example, in cyro-EM data, many 2D images of one molecule that is under molecule dynamic are captured. To reconstruct the 3D structure and understand the dynamic, we should know the “stage” of each image, i.e. t_i . When p is large, the noise is large and the problem becomes very challenging. In this work, we present a spectral method to recover t_i . With manifold learning, we present the consistency of our method, even in the presence of high-dimensional noise. It is the first work to guarantee the consistency of the spectral seriation algorithm for complicated cases under the high dimensional noise. In numerical analysis, we show its power.

DEPARTMENT OF STATISTICS AND DATA SCIENCE, NATIONAL UNIVERSITY OF SINGAPORE
Email address: `wanjie.wang@nus.edu.sg`

THEORY OF FUNCTIONAL PRINCIPAL COMPONENT ANALYSIS FOR DISCRETELY OBSERVED DATA

HANG ZHOU, DONGYI WEI AND FANG YAO

Classification MCS2020: 62R10; 62G20

Keywords: Kernel smoothing; perturbation series; phase transition; optimal convergence

Functional data analysis is an important research field in statistics which treats data as random functions drawn from some infinite-dimensional functional space, and functional principal component analysis (FPCA) based on eigen-decomposition plays a central role for data reduction and representation. After nearly three decades of research, there remains a key problem unsolved, namely, the perturbation analysis of covariance operator for diverging number of eigencomponents obtained from noisy and discretely observed data. This is fundamental for studying models and methods based on FPCA, while there has not been substantial progress since Hall, Müller and Wang (2006)'s result for a fixed number of eigenfunction estimates. In this work, we aim to establish a unified theory for this problem, obtaining upper bounds for eigenfunctions with diverging indices in both the $\mathcal{L}^2$ and supremum norms, and deriving the asymptotic distributions of eigenvalues for a wide range of sampling schemes. Our results provide insight into the phenomenon when the $\mathcal{L}^2$ bound of eigenfunction estimates with diverging indices is minimax optimal as if the curves are fully observed, and reveal the transition of convergence rates from nonparametric to parametric regimes in connection to sparse or dense sampling. We also develop a double truncation technique to handle the uniform convergence of estimated covariance and eigenfunctions. The technical arguments in this work are useful for handling the perturbation series with noisy and discretely observed functional data and can be applied in models or those involving inverse problems based on FPCA as regularization, such as functional linear regression.

Let $X(t)$ be a square integrable stochastic process on $[0, 1]$, and let $X_i(t)$ be independent and identically distributed (i.i.d.) copies of $X(t)$. The mean and covariance functions of $X(t)$ are denoted by $\mu(t) = \mathbb{E}\{X(t)\}$ and $C(s, t) = \mathbb{E}[\{X(s) - \mu(s)\}\{X(t) - \mu(t)\}]$, respectively. According to Mercer's Theorem, $C(s, t)$ has the spectral decomposition

$$(1) \quad C(s, t) = \sum_{k=1}^{\infty} \lambda_k \phi_k(s) \phi_k(t),$$

where $\lambda_1 > \lambda_2 > \dots > 0$ are eigenvalues and $\{\phi_j\}_{j=1}^{\infty}$ are the corresponding eigenfunctions. We further assume the eigenvalues λ_j are decreasing with $j^{-a} \gtrsim \lambda_j \gtrsim \lambda_{j+1} + j^{-a-1}$ for $a > 1$ and each $j \geq 1$.

We usually use n to denote the sample size, which is the number of subjects corresponding to random functions, and N_i to denote the number of observations for

the i th subject. Specifically, the actual observations for each X_i are given by

$$(2) \quad \{(t_{ij}, X_{ij}) \mid X_{ij} = X_i(t_{ij}) + \varepsilon_{ij}, j = 1, \dots, N_i\},$$

where ε_{ij} are random copies of ε , with $\mathbb{E}(\varepsilon) = 0$ and $\text{Var}(\varepsilon) = \sigma_X^2$. We further assume the measurements errors $\{\varepsilon_{ij}\}_{i,j}$ are independent of X_i . The local linear estimator of the covariance function is given by $\hat{C}(s, t) = \hat{\beta}_0$,

$$(3) \quad \begin{aligned} (\hat{\beta}_0, \hat{\beta}_1, \hat{\beta}_2) = \underset{\beta_0, \beta_1, \beta_2}{\text{argmin}} \sum_{i=1}^n v_i \sum_{1 \leq l_1 \neq l_2 \leq N_i} \{ \delta_{ijl} - \beta_0 - \beta_1(t_{il_1} - s) - \beta_2(t_{il_2} - t) \}^2 \\ \times \frac{1}{h} K\left(\frac{t_{il_1} - s}{h}\right) \frac{1}{h} K\left(\frac{t_{il_2} - t}{h}\right), \end{aligned}$$

where K is a symmetric, Lipschitz continuous density kernel on $[-1, 1]$ and h is the tuning parameter.

The estimated covariance function $\hat{C}(s, t)$ can be expressed as an empirical version of the spectral decomposition in (1), given by

$$(4) \quad \hat{C}(s, t) = \sum_{k=1}^{\infty} \hat{\lambda}_k \hat{\phi}_k(s) \hat{\phi}_k(t),$$

where $\hat{\lambda}_k$ and $\hat{\phi}_k$ are the estimated eigenvalues and eigenfunctions, respectively. These estimates are obtained by solving the eigenequation for $\hat{C}(s, t)$:

$$\int \hat{C}(s, t) \hat{\phi}_k(t) dt = \hat{\lambda}_k \hat{\phi}_k(s), \quad \text{with normalization} \quad \int \hat{\phi}_k^2(s) ds = 1.$$

For specificity, we assume $\langle \hat{\phi}_k, \phi_k \rangle \geq 0$.

The following theorem is one of our main results. It characterizes the $\mathcal{L}^2$ convergence of the estimated eigenfunctions with diverging indices for random design case.

Theorem 1. *Under some regularity assumptions, for all $j \leq m \in \mathbb{N}_+$ satisfies $m^{2a+2}/n \rightarrow 0$, $m^{2a+2}/(n\bar{N}_2^2 h^2) \rightarrow 0$ and $h^2 \max\{m^{a+c}, m^{4a} \log n\} \lesssim 1$,*

$$(5) \quad \|\hat{\phi}_j - \phi_j\|^2 = O_P \left(\frac{j^2}{n} \left\{ 1 + \frac{j^{2a}}{\bar{N}_2^2} \right\} + \frac{j^a}{n\bar{N}_2 h} \left(1 + \frac{j^a}{\bar{N}_2} \right) + h^4 j^{2c+2} \right),$$

where $\bar{N}_2 = (n^{-1} \sum_{i=1}^n N_i^{-2})^{-1/2}$.

Similar to the phase transitions of mean and covariance functions studied in Cai and Yuan (2011) and Zhang and Wang (2016), Corollary 2 presents a systematic partition that ranges from “sparse” to “dense” schemes for eigenfunction estimation under the random design, which is meaningful for FPCA-based models and methods.

Corollary 2. *Under some regularity assumptions, for each $j \leq m$ and let*

$$h_{opt}(j) = (n\bar{N}_2)^{-1/5} j^{(a-2c-2)/5} (1 + j^a/\bar{N}_2)^{1/5},$$

(a) *If $\bar{N}_2 \gtrsim j^a$,*

$$\|\hat{\phi}_j - \phi_j\|^2 = O_P \left(\frac{j^2}{n} + \frac{j^{(4a+2c+2)/5}}{(n\bar{N}_2)^{4/5}} \right).$$

In addition, if $\bar{N}_2 \geq n^{1/4} j^{a+c/2-2}$,

$$\|\hat{\phi}_j - \phi_j\|^2 = O_P \left(\frac{j^2}{n} \right).$$

(b) If $\bar{N}_2 = o(j^a)$,

$$\|\hat{\phi}_j - \phi_j\|^2 = O_P \left(\frac{j^{2a+2}}{n\bar{N}_2^2} + \frac{j^{(8a+2c+2)/5}}{(n\bar{N}_2^2)^{4/5}} \right).$$

The following theorem studies the asymptotic distribution of eigenvalues.

Theorem 3. Under some regularity assumptions, for all $j \leq m \in \mathbb{N}_+$ satisfy $m = o(n^{1/(2a+4)})$, $hm^{2a} \log n \lesssim 1$ and $h^4 m^{2a+2c} \lesssim 1$

$$\Sigma_n^{-\frac{1}{2}} \left(\frac{\hat{\lambda}_j - \lambda_j}{\lambda_j} - K_2 h^2 \int \phi_j^{(2)}(u) \phi_j(u) du + o(h^2) \right) \xrightarrow{d} \mathcal{N}(0, 1),$$

where

$$\begin{aligned} \Sigma_n = & \frac{4!P_0(N)}{n} \frac{\mathbb{E}(\xi_j^4) - \lambda_j^2}{\lambda_j^2} + 4! \frac{P_1(N)}{n} \frac{\int \{C(u, u) + \sigma_X^2\} \frac{\phi_j^2(u)}{f(u)} du}{\lambda_j} \\ & + 4 \frac{P_2(N)}{n} \frac{1}{\lambda_j^2} \left(\left[\int \{C(u, u) + \sigma_X^2\} \frac{\phi_j^2(u)}{f(u)} du \right]^2 - \iint C(u, v)^2 \frac{\phi_j^2(u) \phi_j^2(v)}{f(u) f(v)} dudv \right) \end{aligned}$$

with $K_2 = \int u^2 K(u) du$ and

$$P_0(N) = \frac{1}{n} \sum_{i=1}^n \frac{(N_i - 2)(N_i - 3)}{N_i(N_i - 1)}, \quad P_1(N) = \frac{1}{n} \sum_{i=1}^n \frac{(N_i - 2)}{N_i(N_i - 1)}, \quad P_2(N) = \frac{1}{n} \sum_{i=1}^n \frac{1}{N_i(N_i - 1)}.$$

The following theorem gives the uniform convergence for estimated eigenfunctions with diverging indices.

Theorem 4. Under some regularity assumptions, for all $j \leq m$ satisfying $m^{2a+2}/n = o(1)$, $m^{2a+2}/(n\bar{N}_2^2 h^2) = o(1)$, $h^4 m^{2a+2c} \lesssim 1$ and $hm^a \log n \lesssim 1$,

$$\begin{aligned} (6) \quad \|\hat{\phi}_j - \phi_j\|_\infty = & O \left(\frac{j}{\sqrt{n}} (\sqrt{\ln n} + \ln j) \left\{ 1 + \frac{j^a}{\bar{N}_2} + \sqrt{\frac{j^{a-1}}{\bar{N}_2 h}} \left(1 + \sqrt{\frac{j^a}{\bar{N}_2}} \right) \right\} \right. \\ & \left. + j^a \left| \frac{\ln n}{n} \right|^{1-\frac{1}{\alpha}} \left| j^{1/2} + \frac{\ln n}{\bar{N}_2 h} \right|^{1-\frac{1}{\alpha}} h^{-\frac{1}{\alpha}} + h^2 j^{c+1} \log j \right) \text{ a.s.} \end{aligned}$$

Further details can be found in the preprint: Zhou, H., Wei D. and Yao F., 2022. Theory of functional principal component analysis for discretely observed data. arXiv:2209.08768

REFERENCES

- Cai, T Tony T. T. Yuan, Ming M. (2011). Optimal estimation of the mean function based on discretely sampled functional data: Phase transition. Ann. Statist. 39 2330–2355.
- Hall, Peter P., Müller, Hans-Georg H.-G. Wang, Jane-Ling J.-L. (2006). Properties of principal component methods for functional and longitudinal data analysis. Ann. Statist. 1493–1517.
- Zhang, Xiaoke X. Wang, Jane-Ling J.-L. (2016). From sparse to dense functional data and beyond. Ann. Statist. 44 2281–2321.

DEPARTMENT OF STATISTICS, UNIVERSITY OF CALIFORNIA AT DAVIS
SCHOOL OF MATHEMATICAL SCIENCES, PEKING UNIVERSITY
SCHOOL OF MATHEMATICAL SCIENCES AND CENTER FOR STATISTICAL SCIENCE

Email address: hgzhou@ucdavis.edu, jnwdyi@pku.edu.cn and fyao@math.pku.edu.cn

MANIFOLD LEARNING: AN INVITATION TO DATA SCIENCE

ZHIGANG YAO

The manifold fitting problem can go back to H. Whitney's work in the early 1930s (Whitney (1992)), and finally has been answered in recent years by C. Fefferman's works (Fefferman, 2006, 2005). The solution to the Whitney extension problem leads to new insights for data interpolation and inspires the formulation of the Geometric Whitney Problems (Fefferman et al. (2020, 2021a)): Assume that we are given a set $Y \subset \mathbb{R}^D$. When can we construct a smooth d -dimensional submanifold $\widehat{M} \subset \mathbb{R}^D$ to approximate Y , and how well can $\widehat{M}$ estimate Y in terms of distance and smoothness? To address these problems, various mathematical approaches have been proposed (see Fefferman et al. (2016, 2018, 2021b)). However, many of these methods rely on restrictive assumptions, making extending them to efficient and workable algorithms challenging. As the manifold hypothesis (non-Euclidean structure exploration) continues to be a foundational element in data science, the manifold fitting problem, merits further exploration and discussion within the modern data science community.

This talk will be partially based on recent works of Yao and Xia (2019), Yao, Su, Li and Yau (2022), Yao, Su, and Yau (2023) and Yao, Li, Lu, and Yau (2023).

NATIONAL UNIVERSITY OF SINGAPORE

Email address: zhigang.yao@nus.edu.sg

TWO-SAMPLE TESTS FOR EQUAL DISTRIBUTIONS IN SEPARABLE METRIC SPACES: A NEW DISTANCE-BASED APPROACH

JIN-TING ZHANG

Classification AMS 2020:62G10,62H10,62H15.

Keywords: two-sample test, equal distribution, distance-based test, high-dimensional data, functional data

1. INTRODUCTION

With the surge in advanced data collection methods, researchers are increasingly analyzing complex data objects within separable metric spaces across disciplines. Motivated by real-world datasets such as gene expression and economic indicators, we develop a robust distance-based test to evaluate distributional equality in such data. This paper addresses critical issues in existing methods, offering computational efficiency and accuracy for diverse applications; see [6] for more details.

2. METHODOLOGY

The proposed test is formulated using a distance-based statistic that leverages the negative definiteness property of metrics[2,3]. By deriving its asymptotic null distribution as a χ^2 -type mixture, we introduce a rapid three-cumulant matching approximation [4] to bypass the computational cost of permutations that are widely adopted as in [2,3]. The test's asymptotic power and root-n consistency are established under local alternatives. These theoretic properties are not established for the methods developed in [2,3].

3. RESULTS AND SIMULATIONS

Extensive simulations demonstrate the test's superior size control and power across various settings, including high-dimensional and correlated data. Compared to MMD [1] and energy tests [2,3], our method exhibits consistent performance advantages, particularly in computational efficiency. Empirical validation using gene expression data confirms its ability to discern distributional differences effectively.

4. APPLICATIONS

We apply the proposed test to two datasets: (1) high-dimensional gene expression data distinguishing normal and tumor colon tissues, and (2) functional data on the Gini index across countries. The first dataset, with its dimension much larger than its sample size, is available at <http://genomics-pubs.princeton.edu/oncology/affydata/> and the second dataset is downloaded at <https://data.worldbank.org/indicator/SI.POV.GINI>. Results show the proposed test's robustness against data scaling and sensitivity to kernel

parameter choices in competing methods as developed in [1,2,3], highlighting its practicality in diverse contexts.

5. CONCLUSION

This study presents a versatile, efficient, and statistically robust approach for two-sample distribution testing in separable metric spaces. Future work includes extending this framework to multi-sample scenarios and exploring other functional data applications.

REFERENCES

- [1] A. Gretton, K. M. Borgwardt, M. J. Rasch, B. Schölkopf, and A. Smola. A kernel two-sample test. *Journal of Machine Learning Research*, 13:723773, 2012.
- [2] G. J. Székely and M. L. Rizzo. Testing for equal distributions in high dimension. *Inter Stat*, November(5), 2004.
- [3] G. J. Székely and M. L. Rizzo. Energy statistics: A class of statistics based on distances. *Journal of statistical planning and inference*, 143(8):12491272, 2013.
- [4] J.-T. Zhang. Approximate and asymptotic distributions of chi-squared-type mixtures with applications. *Journal of the American Statistical Association*, 100(469):273285, 2005.
- [5] J.-T. Zhang, J. Guo, and B. Zhou. Testing equality of several distributions in separable metric spaces: A maximum mean discrepancy based approach. *Journal of Econometrics*, page 105286, 2022.
- [6] J.-T. Zhang, M. Qian, and T. Zhu. Two-sample tests for equal distributions in separable metric spaces: a new distance-based approach. Unpublished manuscript, 2024.

DEPARTMENT OF STATISTICS AND DATA SCIENCE, NATIONAL UNIVERSITY OF SINGAPORE
E-mail address: stazjt@nus.edu.sg