

Inference on covariance operators via concentration inequalities: k-sample tests, classification, and clustering via Rademacher complexities

Adam B. Kashlak¹, John A. D. Aston², and Richard Nickl²

¹Cambridge Centre for Analysis, University of Cambridge, Wilberforce Rd, Cambridge, CB3 0WB, U.K

²Statistical Laboratory, University of Cambridge, Wilberforce Rd, Cambridge, CB3 0WB, U.K

April 22, 2016

Abstract

We propose a novel approach to the analysis of covariance operators making use of concentration inequalities. First, non-asymptotic confidence sets are constructed for such operators. Then, subsequent applications including a k sample test for equality of covariance, a functional data classifier, and an expectation-maximization style clustering algorithm are derived and tested on both simulated and phoneme data.

1 Introduction

Functional data spans many realms of applications from medical imaging, (Jiang et al., 2014), to speech and linguistics, (Pigoli et al., 2014), to genetics, (Panaretos et al., 2010). General inference techniques for functional data are one area of analysis that has received much attention in recent years from the construction of confidence sets, to other topics such as k -sample tests, classification, and clustering of functional data. Most testing methodology treats the data as continuous L^2 valued functions and subsequently reduces the problem to a finite dimensional one through expansion in some orthogonal basis such as the often utilized Karhunen-Loève expansion (Horváth and Kokoszka, 2012). However, inference in the more general Banach space setting, making use of a variety of norms, has not been addressed adequately. We propose a novel methodology for performing fully functional inference through the application of concentration inequalities; for general concentration of measure results, see Ledoux (2001) and Boucheron et al. (2013). Special emphasis is given to inference on covariance operators, which offers a fruitful way to analyze functional data.

Imagine multiple samples of speech data collected from multiple speakers. Each speaker will have his or her own sample covariance operator taking into account the unique variations of his or her speech and language. An exploratory researcher may want to find natural clusters amidst the speakers perhaps corresponding to gender, language, or regional accent. Meanwhile, a linguist studying the similarities between languages may want to test for the equality of such covariances. A computer scientist may need to implement an algorithm that when given speech data quickly identifies what language is being spoken and furthermore parses the sound clip and identifies each individual phoneme in order to process the speech into text. Our proposed method has the versatility to yield statistical tests that answer these questions as well as others.

Past methods for analyzing covariance operators (Panaretos et al., 2010; Fremdt et al., 2013) rely on the Hilbert-Schmidt setting for their inference. However, the recent work of Pigoli et al. (2014) argues that the use of the Hilbert-Schmidt metric ignores the geometry of the covariance operators and that more statistical power can be gained by using alternative metrics. Hence, we approach such inference for covariance operators in the full generality of Banach spaces, by using a non-asymptotic concentration of measure approach, which has previously been used in nonparametric statistics and machine learning, sometimes under the name of ‘Rademacher complexities’ (Koltchinskii, 2001, 2006; Bartlett et al., 2002; Bartlett and Mendelson, 2003; Giné and Nickl, 2010; Arlot et al., 2010; Lounici and Nickl, 2011; Kerkycharian et al., 2012; Fan, 2011). These concentration inequalities provide a natural way to construct non-asymptotic confidence regions.

2 Definitions and notation

Generally, we will consider functional data to be in the Hilbert space $L^2(I)$ for $I \subset \mathbb{R}$. Our estimated covariance operators of interest are treated as Banach space valued random variables. Let

$$Op(L^2) = \{T : L^2(I) \rightarrow L^2(I) \mid \exists M \geq 0, \|T\phi\|_{L^2} \leq M \|\phi\|_{L^2} \ \forall \phi \in L^2(I)\}$$

denote the space of all bounded linear operators mapping L^2 into L^2 . This is where our covariance operators will live.

The metrics that will be investigated are those that correspond to the p -Schatten norms. When $p \neq 2$, these are not Hilbert norms. Hence, we will consider the general Banach space setting.

Definition 2.1 (p -Schatten Norm) *Given separable Hilbert spaces H_1 and H_2 , a bounded linear operator $\Sigma : H_1 \rightarrow H_2$, and $p \in [1, \infty)$, then the p -Schatten norm is $\|\Sigma\|_p^p = \text{tr}((\Sigma^* \Sigma)^{p/2})$. For $p = \infty$, the Schatten norm is the operator norm: $\|\Sigma\|_\infty = \sup_{f \in H_1} (\|\Sigma f\|_{H_2} / \|f\|_{H_1})$. In the case that Σ is compact, self-adjoint, and trace-class, then given the associated eigenvalues $\{\lambda_i\}_{i=1}^\infty$, the p -Schatten norm coincides with the ℓ^p norm of the eigenvalues:*

$$\|\Sigma\|_p^p = \begin{cases} \|\lambda\|_{\ell^p}^p = \sum_{i=1}^\infty |\lambda_i|^p, & p \in [1, \infty) \\ \max_{i \in \mathbb{N}} |\lambda_i|, & p = \infty \end{cases}$$

In order to construct a covariance operator from a sample of functional data, the notion of tensor product is required. Let $f, g \in L^2(I)$ and ϕ in the dual space $L^2(I)^*$ with inner product $\langle f, \phi \rangle = \phi(f)$. The tensor product, $f \otimes g$, is the rank one operator defined by $(f \otimes g)\phi = \langle g, \phi \rangle f = \phi(g)f$.

Secondly, we will implement a Rademacher symmetrization technique in the concentration inequalities to come. This requires the use of the namesake Rademacher random variables.

Definition 2.2 (Rademacher Distribution) *A random variable $\varepsilon \in \mathbb{R}$ has a Rademacher distribution if $\text{pr}(\varepsilon = 1) = \text{pr}(\varepsilon = -1) = 1/2$.*

One particularly fruitful avenue of functional data analysis is the analysis of covariance operators. Such an approach to functional data has been discussed by Panaretos et al. (2010) for DNA microcircles, by Fremdt et al. (2013) for the egg laying practices of fruit flies, and by Pigoli et al. (2014) with application to differentiating spoken languages.

Definition 2.3 (Covariance Operator) *Let $I \subseteq \mathbb{R}$, and let f be a random function (variable) in $L^2(I)$ with $\mathbb{E}(\|f\|_{L^2}^2) < \infty$ and mean zero. The associated covariance operator $\Sigma_f \in Op(L^2)$ is defined as $\Sigma_f = \mathbb{E}(f \otimes f) = \mathbb{E}(\langle f, \cdot \rangle f)$.*

If $I = \{i_1, \dots, i_m\}$ has finite cardinality, then $f = (f_1, \dots, f_m)$ is a random vector in $\mathbb{R}^m$ and for some fixed vector $v \in \mathbb{R}^m$, $\mathbb{E}(\langle f, v \rangle f) = \mathbb{E}(f f^T) v$ where $\Sigma_f = \mathbb{E}(f f^T)$ is the usual covariance matrix. Covariance operators are integral operators with the kernel function $c_f(s, t) = \text{cov}\{f(s), f(t)\} \in L^2(I \times I)$. Such operators are characterized by the result that for $f \in L^2(I)$, Σ_f is a covariance operator if and only if it is trace-class, self-adjoint, and compact on $L^2(I)$ where the symmetry follows immediately from the definition and the finite trace norm comes from Parseval's equality (Bosq, 2012, Theorem 1.7)(Horváth and Kokoszka, 2012, Section 2.3).

Furthermore, working under the assumption that $\mathbb{E}(\|f\|_{L^2}^4) < \infty$, we will require tensor powers of covariance operators denoted as $\Sigma^{\otimes 2} : Op(L^2) \rightarrow Op(L^2)$. For a basis $\{e_i\}_{i=1}^\infty \in L^2(I)$ with corresponding basis $\{e_i \otimes e_j\}_{i,j=1}^\infty$ for $Op(L^2(I))$, the previous definition is extended to $\Sigma^{\otimes 2} = \langle \Sigma, \cdot \rangle \Sigma$ where for $\Sigma_1, \Sigma_2 \in Op(L^2)$ with $\Sigma_1 = \sum_{i,j=1}^\infty \lambda_{i,j} e_{i,j}$ and $\Sigma_2 = \sum_{i,j=1}^\infty \gamma_{i,j} e_{i,j}$, then $\langle \Sigma_1, \Sigma_2 \rangle = \sum_{i,j} \lambda_{i,j} \gamma_{i,j}$. Specifically for covariance operators, the tensor power takes on a similar integral operator form with kernel $c_{\Sigma}(s, t, u, v) = \text{cov}(f(s), f(t)) \text{cov}(f(u), f(v))$.

Given an Hilbert space H with inner product $\langle \cdot, \cdot \rangle$, the adjoint of a bounded linear operator $\Sigma : H \rightarrow H$, denoted as Σ^* , is the unique operator such that $\langle \Sigma f, g \rangle = \langle f, \Sigma^* g \rangle$ for $f, g \in H$. The existence of which is given by the Riesz representation theorem (Rudin, 1987, chapter 2). For self-adjoint operators, such as the covariance operators of interest, $\Sigma = \Sigma^*$.

We begin with a sample of functional data. Let $f_1, \dots, f_n \in L^2(I)$ be independent and identically distributed observations with mean zero and covariance operator Σ . Let the sample mean be $\bar{f} = n^{-1} \sum_{i=1}^n f_i$ and the empirical estimate of Σ be $\hat{\Sigma} = n^{-1} \sum_{i=1}^n (f_i - \bar{f}) \otimes (f_i - \bar{f})$. The initial goal is to construct a confidence set for Σ with respect to some metric $d(\cdot, \cdot)$ of the form $\{\Sigma : d(\hat{\Sigma}, \Sigma) \leq r(n, \hat{\Sigma}, \alpha)\}$, which has coverage $1 - \alpha$ for any desired $\alpha \in [0, 1]$ and a radius r depending only on the data and α . Such a confidence set can be utilized for a wide variety of statistical analyses.

3 Constructing a confidence set

3.1 Confidence sets for the mean in Banach spaces

The goal of this section is to construct a non-asymptotic confidence region in the Banach space setting. These can subsequently be specialized to our case of interest, covariance operators, when the X_i below are replaced with $f_i^{\otimes 2}$.

The construction of our confidence set is based on Talagrand's concentration inequality (Talagrand, 1996) with explicit constants. This inequality is typically stated for empirical processes (Giné and Nickl, 2016, Theorem 3.3.9 and 3.3.10), but applies to random variables with values in a separable Banach space $(B, \|\cdot\|_B)$ as well by simple duality arguments (Giné and Nickl, 2016, Example 2.1.6). Let $X_1, \dots, X_n \in (B, \|\cdot\|_B)$ be mean zero independent and identically distributed Banach space valued random variables with $\|X_i\|_B \leq U$ for all $i = 1, \dots, n$ where U is some positive constant. Furthermore, let $\langle \cdot, \cdot \rangle : B \times B^* \rightarrow \mathbb{R}$ such that for $X \in B$ and $\phi \in B^*$ then $\langle X, \phi \rangle = \phi(X)$. Define

$$Z = \sup_{\|\phi\|_{B^*} \leq 1} \sum_{i=1}^n \langle X_i, \phi \rangle = \left\| \sum_{i=1}^n X_i \right\|_B, \quad \sigma^2 = \frac{1}{n} \sum_{i=1}^n \sup_{\|\phi\|_{B^*} \leq 1} \mathbb{E} \left(\langle X_i, \phi \rangle^2 \right),$$

where the supremum is taken over a countably dense subset of the unit ball of B^* . Furthermore, define $v = 2UE(Z) + n\sigma^2$. Then, $\text{pr}\{Z > \mathbb{E}(Z) + r\} \leq \exp\{-r^2/(2v + 2rU/3)\}$. Rewriting Z as $n\|\bar{X} - \mathbb{E}(\bar{X})\|_B$ results in

$$\text{pr} \left[\|\bar{X} - \mathbb{E}(\bar{X})\|_B > \mathbb{E}\{\|\bar{X} - \mathbb{E}(\bar{X})\|_B\} + r \right] < \exp \left(\frac{-n^2 r^2}{2v + 2nrU/3} \right)$$

where $\|X_i\|_B < U$ and $v = 2nUE\{\|\bar{X} - \mathbb{E}(\bar{X})\|_B\} + n\sigma^2$.

The above tail bound incorporates the unknown $\mathbb{E}(\|\bar{X} - \mathbb{E}(\bar{X})\|_B)$. Consequently, a symmetrization technique is used. This term is replaced by the norm of the Rademacher average $R_n = n^{-1} \sum_{i=1}^n \varepsilon_i (X_i - \bar{X})$ where the ε_i are independent and identically distributed Rademacher random variables also independent of the X_i . This substitution is justified by invoking the symmetrization inequality (Giné and Nickl, 2016, Theorem 3.1.21),

$$\mathbb{E}(Z) = \mathbb{E} \left\{ \left\| \frac{1}{n} \sum_{i=1}^n (X_i - \mathbb{E}(\bar{X})) \right\|_B \right\} \leq 2\mathbb{E} \left\{ \left\| \frac{1}{n} \sum_{i=1}^n \varepsilon_i (X_i - \bar{X}) \right\|_B \right\} = 2\mathbb{E}(\|R_n\|_B).$$

In practice, the coefficient of 2 is unnecessary as averaging the data already has a symmetrizing effect, and if starting from symmetric data, then $X - \bar{X}$ and $\varepsilon(X - \bar{X})$ are equal in distribution. This result allows us to replace the original expectation with the expectation of the Rademacher average. Furthermore, Talagrand's inequality also applies to R_n . Hence, the Rademacher average concentrates strongly about its expectation, which justifies dropping the expectation. In practice, one can use the intermediary $\mathbb{E}_\varepsilon(\|R_n\|_B)$, which can be approximated for reasonable sized data sets via Monte Carlo simulations of the ε_i . However, this is not strictly necessary, and for large data sets, a single random draw of ε_i will suffice (Giné and Nickl, 2016, Section 3.4.2).

The resulting $(1 - \alpha)$ -confidence set is

$$\left\{ X : \|X - \bar{X}\|_B \leq \|R_n\|_B + \left\{ \frac{2}{n} \log(2\alpha) (\sigma^2 + 2U \|R_n\|_B) \right\}^{1/2} + \frac{U \log(2\alpha)}{3n} \right\}. \quad (3.1)$$

To make use of these results in practice, both the weak variance σ^2 must be estimated for the data and a reasonable choice of U must be made, and a main contribution of this present paper is to propose some theoretically motivated but practically useful non-asymptotic choices for these constants that work for the functional data applications we are investigating. This will be discussed in the setting of covariance operators in the following subsection.

3.2 Confidence sets for covariance operators

To construct a confidence set for covariance operators, let our functional data $f_i \in L^2(I)$ and $f_i^{\otimes 2} = f_i \otimes f_i \in \text{Op}(L^2)$, the Hilbert space of bounded linear operators mapping L^2 to L^2 , such that $(f_i \otimes f_i)\phi = \langle f, \phi \rangle f$ for some $\phi \in L^2$. Hence, for some desired p -Schatten norm, $\|\cdot\|_p$, with $p \in [1, \infty)$ and with conjugate $q = p/(p-1)$, we have

$$Z = \left\| \frac{1}{n} \sum_{i=1}^n f_i \otimes f_i - \mathbb{E}(f_i \otimes f_i) \right\|_p = \left\| \hat{\Sigma} - \Sigma \right\|_p, \quad \sigma^2 = \frac{1}{n} \sum_{i=1}^n \sup_{\|\Pi\|_q \leq 1} \mathbb{E} \left\{ \langle f_i^{\otimes 2} - \mathbb{E}(f_i^{\otimes 2}), \Pi \rangle^2 \right\}$$

for the supremum being taken over a countably dense subset of the unit ball of $\Pi \in Op(L^2)$. The Rademacher average, $R_n = n^{-1} \sum_{i=1}^n \varepsilon_i \{(f_i - \bar{f})^{\otimes 2} - \hat{\Sigma}\}$, is also in $Op(L^2)$, because for any $\phi \in L^2(I)$ and for some $M \in \mathbb{R}$, $\|R_n \phi\|_{L^2} \leq |\varepsilon_i| \|\{(f_i - \bar{f})^{\otimes 2} - \hat{\Sigma}\} \phi\|_{L^2} \leq M \|\phi\|_{L^2}$ since $\{(f_i - \bar{f})^{\otimes 2} - \hat{\Sigma}\}$ is a bounded operator and $|\varepsilon_i| = 1$. Furthermore, in the case that there exists a fixed $c \in \mathbb{R}$ with $\|f_i\|_{L^2} \leq c$ for all i corresponding to a physical bound on the energy of f_i , $\|f_i^{\otimes 2}\|_p = \|f_i\|_{L^2}^2 \leq c^2 = U$. It will be determined below that $U \geq \sigma$ in this case. In general, setting $U = \sigma$ gives good experimental results when f_i is Gaussian as will be discussed in later sections. This results in $v_n \approx \sigma^2/n$. For any $p \in [1, \infty)$ and $\alpha \in [0, 1/2]$, the proposed $(1 - \alpha)$ -confidence set inspired by Equation 3.1 for covariance operators is

$$C_{n,1-\alpha} = \left[\Sigma : \left\| \hat{\Sigma} - \Sigma \right\|_p \leq \|R_n\|_p + \sigma \left\{ -\frac{2}{n} \log(2\alpha) \right\}^{1/2} - \frac{\sigma \log(2\alpha)}{3n} \right].$$

where σ depends on the distribution on the functional data. As a rule of thumb for the choice of σ^2 , as we will now show, is to note that $\sigma^2 \leq \|E(f^{\otimes 4}) - \Sigma^{\otimes 2}\|_p$ and to estimate this bound empirically by $\hat{\sigma}^2$. For example, when the f_i are from a Gaussian process $\hat{\sigma} \leq 2^{1/2} \|\Sigma\|_p$ as explained in detail in Section 3.4.

To calculate the weak variance σ^2 , define $f^{\otimes n} = f \otimes \dots \otimes f$ to be the n -fold tensor product of f with itself and extend the definition of $\langle \cdot, \cdot \rangle : (L^2)^{\otimes 4} \times \{(L^2)^{\otimes 4}\}^* \rightarrow \mathbb{R}$ such that $\langle f^{\otimes 4}, \phi^{\otimes 4} \rangle = \langle f^{\otimes 2}, \phi^{\otimes 2} \rangle^2 = \langle f, \phi \rangle^4 = \phi(f)^4$. For operators $\Pi \in \{(L^2)^{\otimes 2}\}^*$ and $\Xi \in \{(L^2)^{\otimes 4}\}^*$, the weak variance is

$$\begin{aligned} \sigma^2 &= \frac{1}{n} \sum_{i=1}^n \sup_{\|\Pi\|_q \leq 1} E \left\{ \langle f_i^{\otimes 2} - E(f_i^{\otimes 2}), \Pi \rangle^2 \right\} \\ &\leq \frac{1}{n} \sum_{i=1}^n \sup_{\|\Xi\|_q \leq 1} \left\langle E(f_i^{\otimes 4}) - \{E(f_i^{\otimes 2})\}^{\otimes 2}, \Xi \right\rangle \leq \|E(f^{\otimes 4}) - \Sigma^{\otimes 2}\|_p \end{aligned}$$

where the inequality stems from the fact that the supremum is being taken over a larger set. However, in the Hilbert space setting, the dual of the tensor product does coincide with the tensor product of the dual space, and thus the above inequality can be replaced with an equality if the Hilbert-Schmidt norm, 2-Schatten norm, is used. Given a bound $\|f_i\|_{L^2}^2 \leq c^2 = U$, then $\sigma^2 \leq \|E(f^{\otimes 4})\|_p \leq E(\|f\|_{L^2}^4) \leq c^4 = U^2$.

3.3 The weak variance for $p = \infty$

Let E be a countable dense subset of the unit ball of $L^2(I)$. In the case $p = \infty$, we cannot use duality, but can still write Z and σ^2 as suprema over the countable set and achieve the same results as above.

$$\begin{aligned} Z &= \frac{1}{n} \sup_{e \in E} \sum_{i=1}^n \langle \{f_i^{\otimes 2} - E(f_i^{\otimes 2})\} e, e \rangle = \sup_{e \in E} \langle (\hat{\Sigma} - \Sigma)e, e \rangle = \|\hat{\Sigma} - \Sigma\|_\infty, \\ \sigma^2 &= \frac{1}{n} \sum_{i=1}^n \sup_{e_1 \in E} E \left\{ \langle (f_i^{\otimes 2} - \Sigma)e_1, e_1 \rangle^2 \right\} \leq \frac{1}{n} \sum_{i=1}^n \sup_{e_1, e_2 \in E} E \left\{ \langle f_i^{\otimes 2} - \Sigma, e_1 \otimes e_2 \rangle^2 \right\} \\ &\leq \frac{1}{n} \sum_{i=1}^n \sup_{e_1, e_2 \in E} \langle \{E(f_i^{\otimes 4}) - \Sigma^{\otimes 2}\} (e_1 \otimes e_2), e_1 \otimes e_2 \rangle = \|E(f_i^{\otimes 4}) - \Sigma^{\otimes 2}\|_\infty. \end{aligned}$$

As before, if $\|f_i^{\otimes 2}\|_\infty = \|f_i\|_{L^2}^2 \leq c^2 = U$, then $\sigma^2 \leq U^2$.

3.4 The weak variance for Gaussian data

Similarly to the bounded case, we estimate $\|E(f^{\otimes 4}) - \Sigma^{\otimes 2}\|_p$ for Gaussian data. Consider f from a Gaussian process with mean zero and covariance Σ and define $f_s = f(s)$. Strictly speaking these variables are not norm bounded, but similar to concentration results for Gaussian random variables in $\mathbb{R}^d$ (Giné and Nickl, 2016, Theorem 3.1.9), we found below that our methods still work well. The integral kernel can be written as (Isserlis, 1918)

$$\begin{aligned} E(f_s f_t f_u f_v) &= E(f_s f_t) E(f_u f_v) + E(f_s f_u) E(f_t f_v) + E(f_s f_v) E(f_t f_u) \\ &= c_f(s, t) c_f(u, v) + c_f(s, u) c_f(t, v) + c_f(s, v) c_f(t, u). \end{aligned}$$

Hence, we have that $E(f_s f_t f_u f_v) - \Sigma_{s,t} \Sigma_{u,v} = \Sigma_{s,u} \Sigma_{t,v} + \Sigma_{s,v} \Sigma_{t,u}$ and that the operator $E(f^{\otimes 4}) - \Sigma^{\otimes 2}$, which can be thought of as an Hilbert-Schmidt operator on the space $Op(L^2)$, can be represented by the integral kernel $c_f(s, u)c_f(t, v) + c_f(s, v)c_f(t, u)$. These two terms are merely relabeled versions of $\Sigma^{\otimes 2}$. Consequently, using the subadditivity of the norm, $\|E(f^{\otimes 4}) - \Sigma^{\otimes 2}\|_p \leq \|\Sigma^{\otimes 2}\|_p + \|\Sigma^{\otimes 2}\|_p = 2\|\Sigma^{\otimes 2}\|_p$. For example, for the Hilbert-Schmidt norm,

$$\begin{aligned} \|E(f^{\otimes 4}) - \Sigma^{\otimes 2}\|_{HS}^2 &= \iiint \{c_f(s, u)c_f(t, v) + c_f(s, v)c_f(t, u)\}^2 ds dt du dv \\ &= 2\|\Sigma\|_{HS}^4 + 2\iiint c_f(s, u)c_f(s, v)c_f(t, v)c_f(t, u) ds dt du dv \leq 4\|\Sigma\|_{HS}^4. \end{aligned}$$

Lemma 5.1 of Horváth and Kokoszka (2012) gives an explicit form of a covariance operator of Σ in terms of the eigenfunctions of Σ for Gaussian data in the Hilbert-Schmidt setting.

Given λ_i , the eigenvalues of Σ , the spectrum of $\Sigma^{\otimes 2}$ is $\{\lambda_i \lambda_j\}_{i,j=1}^\infty$. Hence, for any of the p -Schatten norms, $\|\Sigma \otimes \Sigma\|_p = \|\Sigma\|_p^2$. Note that in the above calculations, the weak variance depends on the unknown Σ . In practice, this can be replaced by the empirical estimate $\hat{\Sigma}$.

4 Applications

4.1 k Sample Comparison

Testing for the equality of means among multiple sets of data is a common task in data analysis. In the functional setting, there has been recent work on performing such a test on covariance operators in order to test whether or not k sets of curves have similar variation. Panaretos et al. (2010) propose such a method for a two sample test on covariance operators given data from Gaussian processes. Similarly, Fremdt et al. (2013) propose a non-parametric two sample test on covariance operators. Both of these approaches make use of the Karhunen-Loève expansion and, hence, the underlying Hilbert space geometry. Pigoli et al. (2014) take a comparative look at a variety of metrics to rank their statistical power when used in a two sample permutation test.

Following from the results of Pigoli et al. (2014), our method uses the p -Schatten norms with the concentration inequality based confidence sets of the previous section to compare covariance operators. In the two sample setting, we are able to achieve similar statistical power to that of the permutation test after proper tuning of the coefficients in the inequalities. Furthermore, the analytic nature of the concentration approach leads to a significant reduction in computing time, which offers an even more significant savings for larger values of k .

From the confidence set constructed in the previous section, we can devise a test for comparing the empirical covariance operators generated from k samples of functional data. Let the k samples be $f_1^{(1)}, \dots, f_{n_1}^{(1)}, \dots, f_1^{(k)}, \dots, f_{n_k}^{(k)}$ where for each sample i and all elements $j = 1, \dots, n_i$, $f_j^{(i)}$ has covariance $\Sigma^{(i)}$. Our goal is to design a test for the following two hypotheses:

$$H_0 : \Sigma^{(1)} = \dots = \Sigma^{(k)} \quad H_1 : \exists i, j \text{ s.t. } \Sigma^{(i)} \neq \Sigma^{(j)}.$$

To achieve this, a pooled estimate of the weak variance is computed as a weighted average of each sample's individual weak variance in similar style to that of a standard t-test. Let the total data size be $N = n_1 + \dots + n_k$ and σ_i^2 be the weak variance for sample i , then the pooled variance is defined as $\sigma_{\text{pool}}^2 = N^{-1} \sum_{i=1}^k n_i \sigma_i^2$. Given Gaussian data and the p -Schatten norm, for example, this reduces to $\sigma_{\text{pool}}^2 = 2N^{-1} \sum_{i=1}^k n_i \|\Sigma^{(i)}\|_p^2$. In practice, σ_{pool}^2 is estimated from the data for the following confidence regions in order to have those regions only depend on the data.

Taking inspiration from the standard analysis of variance (Casella and Berger, 2002, Chapter 11), let $\hat{\Sigma}^{(i)}$ be the empirical estimate of the covariance operator for the i th sample, and let $\hat{\Sigma}$ be the estimate of the covariance operator for the total data set. Making use of the confidence sets for covariance operators from Section 3.2 gives the rejection region

$$\begin{aligned} \mathcal{C} = \left\{ f : \sum_{i=1}^k \left\| \hat{\Sigma}^{(i)} - \hat{\Sigma} \right\|_p > \sum_{i=1}^k \left\| \sum_{j=1}^{n_i} \varepsilon_{i,j} \left(f_j^{(i)\otimes 2} - \hat{\Sigma} \right) \right\|_p + \right. \\ \left. + \left(\sum_{i=1}^k \frac{\sigma_{\text{pool}}^2}{n_i} \right)^{1/2} (-2 \log 2\alpha)^{1/2} + \left(\sum_{i=1}^k \frac{\sigma_{\text{pool}}}{n_i} \right) \frac{\log 2\alpha}{3} \right\}, \end{aligned}$$

which under the null hypothesis will have size no greater than the desired α .

The size of the test induced by this rejection region is significantly less than the target size α due to the use of multiple concentration inequalities. Hence, tuning the inequalities is required to yield a useful test. Many experiments were run on simulated data sets generated as samples from a Gaussian process with randomly generated covariance operators whose eigenvalues were chosen to decay at a variety of rates. In this setting, the coefficients of $1 - k^{-1/2}$ for the Rademacher term and $(k + 2)/(k + 3)$ for the deviation term were determined experimentally to improve the size of the confidence region:

$$\mathcal{C} = \left[f : \sum_{i=1}^k \left\| \hat{\Sigma}^{(i)} - \hat{\Sigma} \right\|_p > \left(1 - k^{-1/2} \right) \sum_{i=1}^k \left\| \sum_{j=1}^{n_i} \varepsilon_{i,j} \left(f_j^{(i) \otimes 2} - \hat{\Sigma} \right) \right\|_p + \left(\frac{k+2}{k+3} \right) \left\{ \left(\sum_{i=1}^k \frac{\sigma_{\text{pool}}^2}{n_i} \right)^{1/2} (-2 \log 2\alpha)^{1/2} + \left(\sum_{i=1}^k \frac{\sigma_{\text{pool}}}{n_i} \right) \frac{\log 2\alpha}{3} \right\} \right]. \quad (4.1)$$

4.2 Classification of Operators

Classification of functional data has been an area of heavy research over the last two decades. James and Hastie (2001) extend linear discriminant analysis to functional data. Hall et al. (2001) and Glendinning and Herbert (2003) classify with principle components. Ferraty and Vieu (2003) implement kernel estimators. General linear models for functional data are discussed by Müller and Stadtmüller (2005). Delaigle and Hall (2012) analyze the asymptotic properties of the centroid based classifier. Wavelet based classification is detailed by Berlinet et al. (2008) and Chang et al. (2014).

One application of our method beyond classification of functional data is the classification of covariance operators. In the setting of speech analysis, consider multiple speakers and multiple samples of speech from each speaker. The speech samples can be combined into a single sample covariance operator for each speaker. Then, our method can be employed, for example, to classify the covariance operators by speaker gender or speaker language. Evidence that this is a fruitful approach can be found in the analysis of Pigoli et al. (2014) and Pigoli et al. (2015) where a variety of metrics are compared for their efficacy when performing inference on covariance operators. These articles detail the discrepancy between sample covariance operators produced by speakers of different romance languages.

Given k possible labels and n samples of labeled data (Y_i, f_i) with label $Y_i \in \{1, \dots, k\}$ and observation $f_i \in L^2(I)$, our goal is to determine the probability that a newly observed $g \in L^2(I)$ belongs to label $Y = j$. Given such a g , the Bayes classifier chooses the label $y = \arg \max_j \text{pr}(Y = j \mid g)$ where $\text{pr}(Y = j \mid g) = \text{pr}(g \mid Y = j) \text{pr}(Y = j) / \text{pr}(g)$.

Beginning with a training set of n samples with n_j samples of label j , the sample mean of each category is computed: $\bar{f}_j = n_j^{-1} \sum_{i: Y_i=j} f_i$. The probability $\text{pr}(g \mid Y = j)$ above is replaced with $\text{pr}[\|f_j - g\|_{L^2} > E\{\|f_j - E(\bar{f}_j)\|_{L^2}\} + r]$ with the goal of making a decision based on how much more $\bar{f}_j$ differs from g than $\bar{f}_j$ differs from its expectation $E(\bar{f}_j)$. Similar techniques to those in Section 3 as used. Define the Rademacher sum, R_j , and the empirical weak variance, $\hat{\sigma}_j^2$, for label j to be, respectively,

$$R_j = \frac{1}{n_j} \sum_{i: Y_i=j} \varepsilon_i (f_i - \bar{f}_j), \quad \hat{\sigma}_j^2 = \left\| \frac{1}{n_j} \sum_{i: Y_i=j} f_i^{\otimes 2} - \bar{f}_j^{\otimes 2} \right\|_p$$

where ε_i are independent and identically distributed Rademacher random variables. The tail bound for the above probability is then

$$\text{pr}(\|\bar{f}_j - g\|_{L^2} - \|R_j\|_{L^2} > r) < \exp \left(\frac{-n_j r^2}{4 \|R_j\|_{L^2} U + 2 \hat{\sigma}_j^2 + 2rU/3} \right), \quad (4.2)$$

where U is an upper bound on $\|f_i\|_{L^2}$. However, this can be approximated by the Gaussian tail $\exp(-n_j r^2 / 2 \hat{\sigma}_j^2)$. In the simulations of Section 5.3, this approximation actually achieves a better correct classification rate on both Gaussian and t-distributed data. This specifically works on t-distributed data as the estimate in Equation 4.3 below is merely concerned with comparing the tail bounds rather than their specific values. Consequently, the tail for every category is underestimated in the t case, but the ratio remains valid for comparison purposes.

Assuming uniform priors on the labels, the estimate for the probability expression in the Bayes classifier is

$$\text{pr}(Y = j \mid g) \approx \frac{\phi_j(g)}{\sum_{l=1}^k \phi_l(g)}, \quad \phi_j(g) = \exp \left\{ -\frac{n_j}{2} \left(\frac{\|\bar{f}_j - g\|_{L^2} - \|R_j\|_{L^2}}{\hat{\sigma}_j} \right)^2 \right\}. \quad (4.3)$$

This can be extended to the case where an unlabeled observation is a collection of curves $g_1, \dots, g_m$ by replacing $\|f_j - g\|_{L^2}$ in the above expression with $\|\hat{\Sigma}_j - \hat{\Sigma}_g\|_p$ where $\hat{\Sigma}_j$ is the sample covariance of the f_i with label j and $\hat{\Sigma}_g$ is the sample covariance of the g_i . The Rademacher and weak variance terms would also be updated accordingly.

4.3 Clustering of Operator Mixtures

Closely related to the problem of classification is the problem of clustering. Given a sample of functional data, we want to assign one of a finite collection of labels to each curve. For example, in speech processing, one may want to cluster sound clips based the language of the speaker, or, to be discussed in Section 5.4, one may want to separate unlabeled phoneme curves into clusters.

There have been many recently proposed methods for clustering functional data. Many approaches begin by constructing a low dimensional representation of the data in some basis such as modelling the data with a B-spline basis followed by clustering the spline representations with k-means (Abraham et al., 2003). A similar approach makes use of the eigenfunctions of the covariance operator instead of B-splines (Peng and Müller, 2008). In contrast, we will attempt to cluster functions or operators directly via a concentration of measure approach similar to the previously described classification procedure.

Consider the same setting to the previous section of multiple observations from multiple categories. However, now the category labels are missing. This is a functional mixture model where each observed functional datum is a stochastic process with one of k possible covariance operators. In the below experiments, the data will be simulated from a Gaussian process. The goal is to correctly separate the data into k sets. To achieve this, an expectation-maximization style algorithm is implemented.

Let the observed operator data be $S_1, \dots, S_n \in Op(L^2)$ where each $S_i = \text{cov}(f_1^{(i)}, \dots, f_{m_i}^{(i)})$ is a rank m_i operator produced from m_i functional observations. Let the latent label variables be $Y_1, \dots, Y_n \in \{1, \dots, k\}$. Assuming no prior knowledge on the proportions of data in each category, the algorithm is initialized with the Jeffreys prior for the Dirichlet distribution by randomly generating $\rho_{i,\cdot}^{(0)} \sim \text{Dirichlet}(1/2, \dots, 1/2)$, the initial probability vector that $\text{pr}(Y_i = * | f_i)$.

Assuming t iterations of the algorithm have completed, we have a label probability vector $\rho_{i,\cdot}^{(t)}$ for each of the n observations. Given this collection of vectors, the expected proportions of each category can be estimated as $\tau_j^{(t+1)} = n^{-1} \sum_{i=1}^n \mathbb{E}(\mathbf{1}_{Y_i=j}) = n^{-1} \sum_{i=1}^n \rho_{i,j}^{(t)}$. Similarly, a weighted sum of the data, $\hat{\Sigma}_j^{(t+1)}$, and a weighted Rademacher sum, $R_j^{(t+1)}$, can be used to update the estimated covariance operators for each label j :

$$\hat{\Sigma}_j^{(t+1)} = \frac{\sum_{i=1}^n \rho_{i,j}^{(t)} S_i}{\sum_{i=1}^n \rho_{i,j}^{(t)}}, \quad R_j^{(t+1)} = \frac{\sum_{i=1}^n \rho_{i,j}^{(t)} \varepsilon_i (S_i - \hat{\Sigma}_j^{(t+1)})}{\sum_{i=1}^n \rho_{i,j}^{(t)}}.$$

Lastly, a pooled weak variance is required, which is used in place of each individual category weak variance. Otherwise, in practice, one single category captures all of the data points. By defining the pooled covariance operator as $\hat{\Sigma}_{\text{pool}}^{(t+1)} = \sum_{j=1}^k \tau_j^{(t+1)} \hat{\Sigma}_j^{(t+1)}$, then the pooled weak variance in the Gaussian case, for example, is estimated by $2\|\hat{\Sigma}_{\text{pool}}^{(t+1)}\|_p$.

As a result, the label probability vectors $\rho_{i,\cdot}^{(t)}$ can be updated given the $t + 1$ st collection of estimated covariance operators, Rademacher sums, and the pooled covariance operator. From the previous section, Equation 4.3 can be used to determine $\rho_{i,j}^{(t+1)} = \text{pr}(Y_i = j | S_i, \hat{\Sigma}_1^{(t+1)}, \dots, \hat{\Sigma}_k^{(t+1)})$, the probability that observation i belongs to the j th category. This process can be iterated until a local optimum is reached.

5 Numerical Experiments

5.1 Simulated and phoneme data

To test each of the above three applications, experiments were first run on simulated data. These data sets were generated as observations from Gaussian or t-distributed processes. The covaraiance operators were randomly chosen given a specific decay rate for the eigenvalues.

Secondly, the phoneme data to be tested (Ferraty and Vieu, 2003; Hastie et al., 1995) is a collection of 400 log-periodograms for each of five different phonemes: /a/ as in the vowel of “dark”; /ɔ/ as in the first vowel of “water”; /d/ as in the plosive of “dark”; /i/ as in the vowel of “she”; /ʃ/ as in the fricative of “she”. Each curve contains the first 150 frequencies from a 32 ms sound clip sampled at a rate of 16-kHz.

5.2 k Sample Comparison

The above confidence set in Equation 4.1 comparing k samples can be used to refute the null hypothesis that all covariance operators are equal. A two sample permutation test was performed in Pigoli et al. (2014). Given two samples of functional data, $f_1^{(1)}, \dots, f_n^{(1)}$ and $f_1^{(2)}, \dots, f_m^{(2)}$ with associated covariance operators $\Sigma^{(1)}$ and $\Sigma^{(2)}$, respectively, the desired hypotheses to test are

$$H_0 : \Sigma^{(1)} = \Sigma^{(2)} \quad H_1 : \Sigma^{(1)} \neq \Sigma^{(2)}.$$

When using a permutation test, the labels are randomly reassigned M times, and each time, the distance between the two new covariance operators is computed. For sufficiently large M , this procedure will return the exact significance level of the observations with respect to the data set.

A power analysis was performed between the permutation method and our proposed concentration approach using Equation 4.1. Given two different operators $\Sigma^{(1)}$ and $\Sigma^{(2)}$ and $\gamma \geq 0$, an interpolation between the two operators is constructed as $\Sigma^{(\gamma)} = [(\Sigma^{(1)})^{1/2} + \gamma\{S(\Sigma^{(2)})^{1/2} - (\Sigma^{(1)})^{1/2}\}][(\Sigma^{(1)})^{1/2} + \gamma\{S(\Sigma^{(2)})^{1/2} - (\Sigma^{(1)})^{1/2}\}]^*$, where S is an operator minimizing the Procrustes distance, between $\Sigma^{(1)}$ and $\Sigma^{(2)}$, which is $d_{\text{Proc}}(\Sigma^{(1)}, \Sigma^{(2)})^2 = \inf_{S \in U\{L^2(I)\}} \|R^{(1)} - R^{(2)}S\|_2^2$ where $\Sigma^{(i)} = (R^{(i)})(R^{(i)})^*$ and $U\{L^2(I)\}$ is the space of unitary operators on $L^2(I)$ (Pigoli et al., 2014).

Monte Carlo simulations were run in order to estimate the power of each test. Two operators $\Sigma^{(1)}$ and $\Sigma^{(2)}$ with similar eigenvalue decay were compared. For sample size $n = 50$, $\gamma \in \{0, .1, .2, .3, .4, .5\}$, and for sample size $n = 500$, $\gamma \in \{0, .05, .1, .15, .2, .25\}$. For each γ , 5000 samples of size n were generated for $\Sigma^{(1)}$ and $\Sigma^{(\gamma)}$. Equation 4.1 and the permutation method (Pigoli et al., 2014) were both implemented to estimate the empirical power.

Figures 1 and 2 display the results for operators whose eigenvalues decay at a quartic and quadratic rate, respectively. The solid lines indicate the power of the permutation test, and the circle lines indicate the power of our concentration approach. The colors red, green, and blue correspond to the three norms trace, Hilbert-Schmidt, and operator, respectively.

In most cases, the concentration approach is able to achieve the same power to reject the null as does the permutation test. The notable exception is for the trace norm when the eigenvalues decay slowly. The added benefit to the concentration approach is the speed with which it executes. Across all of the Monte Carlo simulations, our concentration approach ran on average 28.14 times faster than the permutation method. This was computed by tracking the amount of computation time each method spent while producing the plots in Figures 1 and 2, which corresponds to 5 values of γ , 2 values of α , 3 different norms, 2 sample sizes, and 5000 replications each resulting in 300,000 function calls for both the permutation and concentration methods. Unlike the other norms, the Hilbert-Schmidt norm can be calculated without explicit computation of the eigenvalues. For each evaluation of the permutation test, 100 permutations of the data were generated, which corresponds to 100 random draws and 100 eigenvalue computations. More accuracy would require even more permutations. In comparison, our concentration approach requires only $3k$ eigenvalue computations and no random draws and hence is only dependent on the number of samples regardless of data size or α .

The proposed k -sample test was also used to compare samples of log-periodogram curves from the spoken phonemes /a/ and /ɔ/. As one can imagine, these vowels can be hard to distinguish; see Section 5.4 for further evidence of this. For $k \in \{2, 3, 4, 5, 6\}$, $k - 1$ disjoint sets of 40 /a/ curves and one set of 40 /ɔ/ curves were randomly sampled from the data set. This was replicated 500 times, and each time Equation 4.1 was used to decide whether or not the k covariance operators were equivalent at the $\alpha = 0.05$ level. The resulting estimated statistical power for each k is

k	2	3	4	5	6
Power	0.00	0.018	0.228	0.656	0.936

In the null setting, the above experiment was rerun except that every disjoint set of curves came from the /a/ set. The resulting experimentally computed test sizes are

k	2	3	4	5	6
Size	0.00	0.00	0.00	0.004	0.072

5.3 Binary and trinary classification

Our concentration of measure (CoM) method is implemented on covariance operators making use of the trace norm $\|\cdot\|_{\text{tr}}$ where for a covariance operator Σ with eigenvalues $\{\lambda_i\}_{i=1}^{\infty}$, $\|\Sigma\|_{\text{tr}} = \sum_{i=1}^n |\lambda_i|$. The trace norm was chosen based on the analysis of the preceding section as well as that of Pigoli et al. (2014) where it achieved the best performance when compared with the other p-Schatten norms. The CoM approach to classification of operators is tested in a variety of simulations against other standard approaches to functional classification. The methods used for comparison are k -nearest

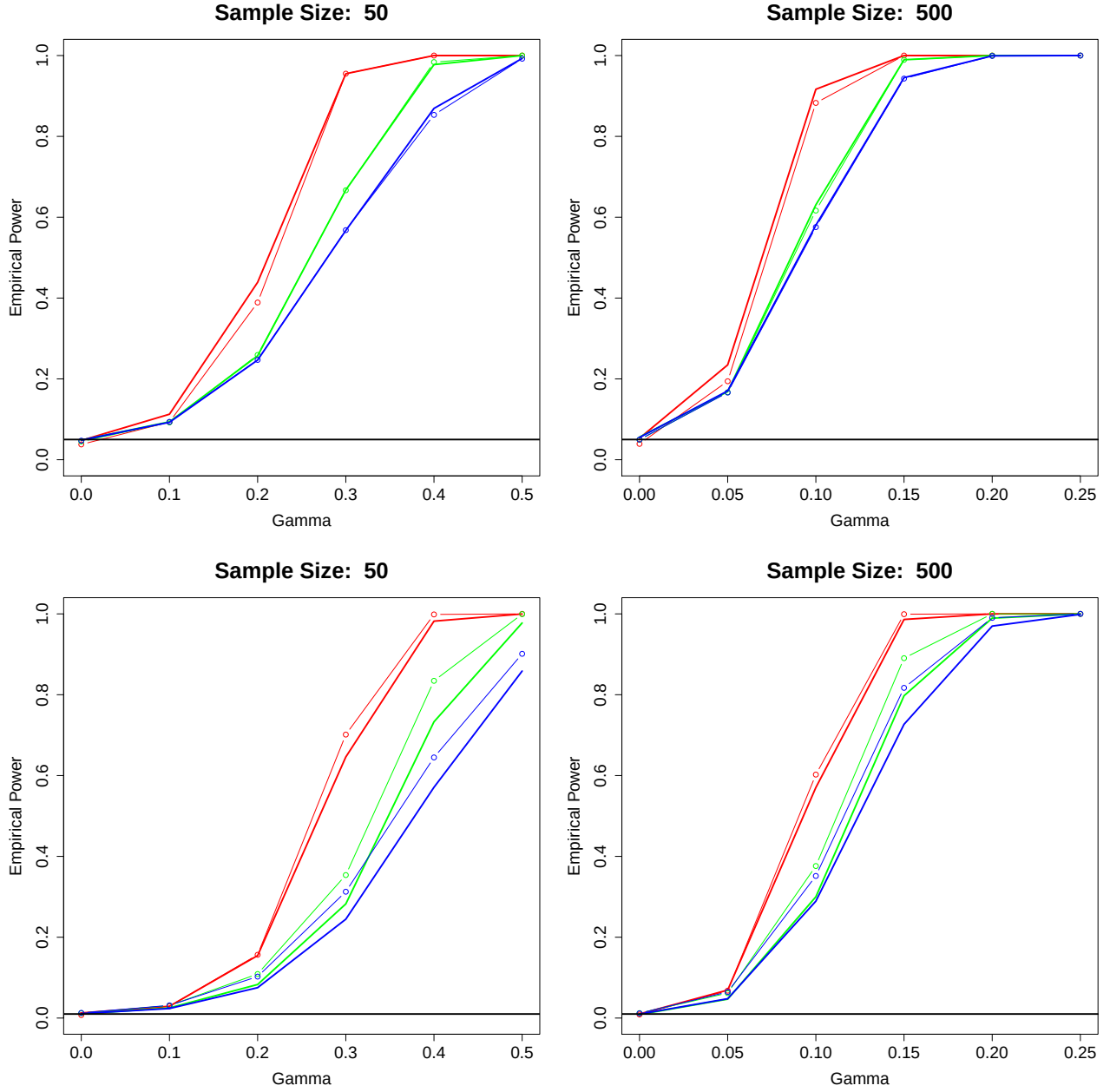


Figure 1: A power analysis for testing whether or not operator $\Sigma^{(1)} = \Sigma^{(\gamma)}$ comparing the permutation method (solid lines) with the concentration approach (circle lines). The size $\alpha = 0.05$ in the top row, and $\alpha = 0.01$ in the bottom. The eigenvalues of the operators decay at a rate $O(k^{-4})$. The red, green, and blue lines respectively correspond to the trace class, Hilbert-Schmidt, and operator norms.

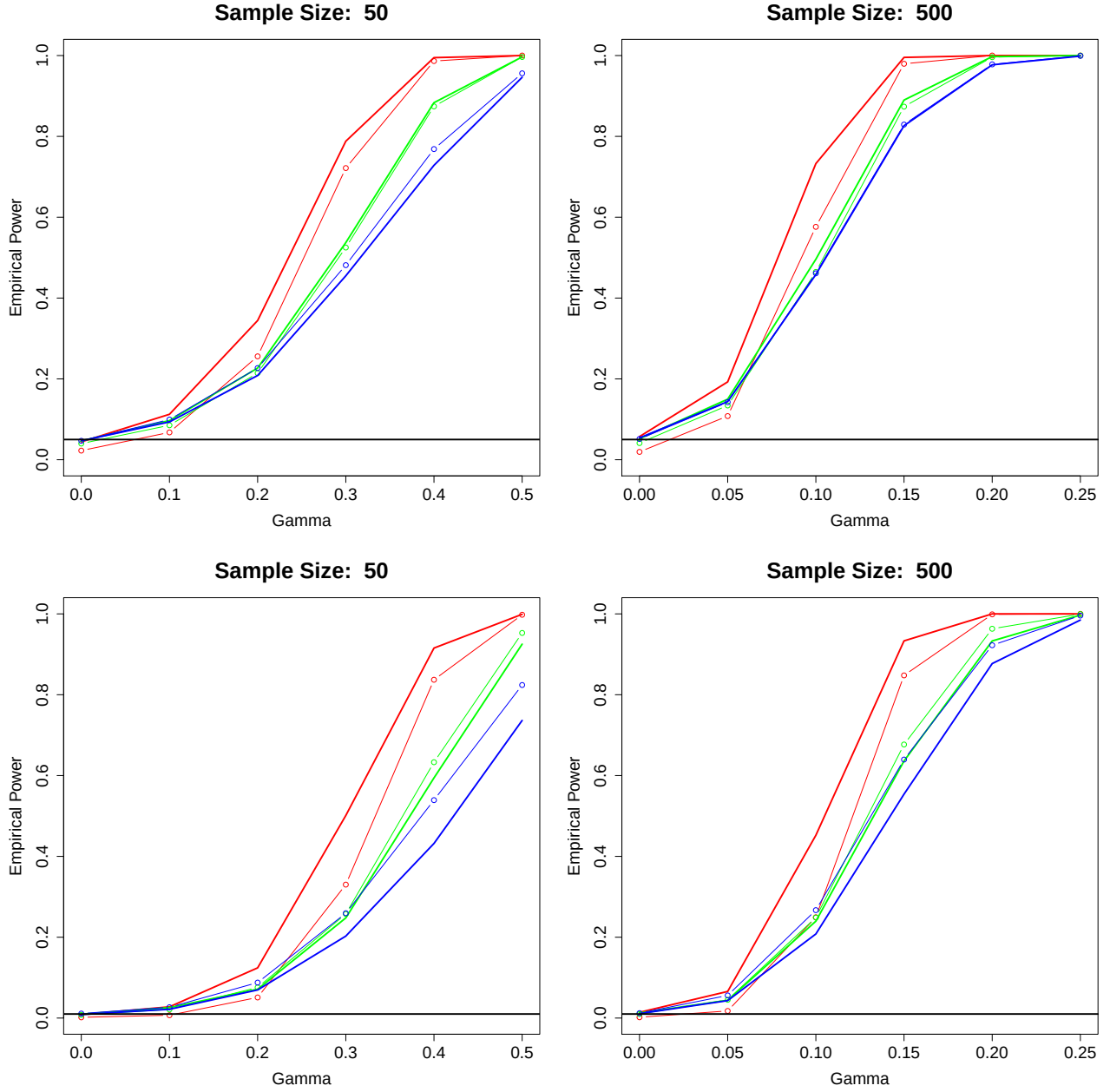


Figure 2: A power analysis for testing whether or not operator $\Sigma^{(1)} = \Sigma^{(\gamma)}$ comparing the permutation method (solid lines) with the concentration approach (circle lines). The size $\alpha = 0.05$ in the top row, and $\alpha = 0.01$ in the bottom. The eigenvalues of the operators decay at a rate $O(k^{-2})$. The red, green, and blue lines respectively correspond to the trace class, Hilbert-Schmidt, and operator norms.

Table 1: A comparison of the performances of five classification methods. The first entry for each method corresponds to classifying the covariance operators as functions. The prime entry corresponds to classifying curves with a majority vote. The estimated percent of correct classification is listed in the table with the sample standard deviation in brackets. The top block comes from Gaussian process data, and the bottom comes from t-process data with 4 degrees of freedom. The highest percentage of each column is marked in bold.

k	1	2	4	8	16
CoM	0.62 (0.05)	0.62 (0.05)	0.76 (0.08)	0.87 (0.06)	0.96 (0.03)
KNN	0.52 (0.04)	0.44 (0.04)	0.57 (0.04)	0.76 (0.04)	0.91 (0.02)
KNN'	.	0.47 (0.03)	0.59 (0.04)	0.74 (0.03)	0.89 (0.02)
Kernel	0.65 (0.04)	0.64 (0.05)	0.75 (0.03)	0.87 (0.03)	0.96 (0.02)
Kernel'	.	0.70 (0.04)	0.82 (0.02)	0.92 (0.02)	0.99 (0.01)
GLM	0.51 (0.04)	0.62 (0.04)	0.74 (0.04)	0.86 (0.02)	0.94 (0.02)
GLM'	.	0.50 (0.04)	0.50 (0.03)	0.50 (0.04)	0.50 (0.04)
Tree	0.57 (0.04)	0.54 (0.04)	0.59 (0.03)	0.66 (0.04)	0.75 (0.03)
Tree'	.	0.55 (0.04)	0.59 (0.04)	0.60 (0.05)	0.60 (0.09)
CoM	0.59 (0.05)	0.62 (0.05)	0.75 (0.07)	0.86 (0.06)	0.95 (0.04)
KNN	0.42 (0.04)	0.45 (0.04)	0.58 (0.04)	0.76 (0.03)	0.92 (0.02)
KNN'	.	0.45 (0.04)	0.58 (0.04)	0.72 (0.03)	0.89 (0.03)
Kernel	0.65 (0.04)	0.64 (0.05)	0.75 (0.04)	0.87 (0.02)	0.96 (0.02)
Kernel'	.	0.68 (0.04)	0.80 (0.03)	0.92 (0.02)	0.99 (0.01)
GLM	0.50 (0.03)	0.62 (0.03)	0.74 (0.04)	0.86 (0.03)	0.94 (0.02)
GLM'	.	0.50 (0.03)	0.50 (0.03)	0.51 (0.04)	0.50 (0.03)
Tree	0.54 (0.04)	0.54 (0.04)	0.59 (0.03)	0.67 (0.03)	0.75 (0.03)
Tree'	.	0.54 (0.04)	0.57 (0.04)	0.59 (0.05)	0.57 (0.06)

CoM, concentration of measure; KNN, k-nearest-neighbours; GLM, generalized linear model; Tree, regression tree.

neighbours (Ferraty and Vieu, 2006), classification using kernel estimators (Ferraty and Vieu, 2003), general linear model (Müller and Stadtmüller, 2005), and regression trees.

The first simulation asks each method to classify observed mean zero Gaussian process data or mean zero t-process data with 4 degrees of freedom. The two covariance operators in question, Σ_1 and Σ_2 , are the sample covariances of the male and of the females of the Berkeley growth curve data (Ramsay and Silverman, 2005). In particular, n collections of k curves were generated from each of Σ_1 and Σ_2 as a training set, and m collections of k curves were generated as a test set. The CoM method was trained on the set of n sample covariances and used to classify each of the m test covariances. The remaining classification methods were trained and tested in two separate ways: By treating each sample covariance as a function and classifying as usual, and by training on all $n \times k$ observations and testing each of the m collections by classifying each constituent curve individually and taking a majority vote with ties settled by a uniform random draw.

For group sizes $k = 1, 2, 4, 8, 16$, 100 simulations were run with $n = 100$ sets of k training curves. To compare the accuracy of each approach $m = 100$ sets of k testing curves were generated for each operator. The accuracy of each method is tabulated in Table 1.

The concentration method performed well against the alternatives. Its performance was on par with the kernel method applied to each covariance operator as a function. Our method was only consistently outperformed by the kernel method implementing the majority vote approach. However, the two operators in question have very similar weak variances. The next simulation demonstrates how the concentration method adapts naturally when the variances of each label significantly differ.

Continuing from the previous simulation, a third operator is constructed from Σ_1 and Σ_2 by averaging these two and then scaling up the non-principle eigenvalues by a factor of 5. This, in some sense, creates a third operator between the first two, but with higher variance. The simulation is carried out precisely as before, but incorporating all three operators. In this setting, our concentration approach demonstrates the best performance. The results are listed in Table 2.

These five methods tested on simulated data were also tested against phoneme data. Across 50 iterations, each set of 400 curves was partitioned at random into an 100 curve training set and a 300 curve testing set. The five classifiers were trained and run on each of the 300×5 curves individually. For our concentration of measure approach, the rank one operator associated to each individual curve was compared with the covariance operator formed from the 100×5

Table 2: A comparison of the performances of five classification methods as in Table 1, but with three potential classes from which to choose. The estimated percent of correct classification is listed in the table with the sample standard deviation in brackets. The top block comes from Gaussian process data, and the bottom comes from t-process data with 4 degrees of freedom. The highest percentage of each column is marked in bold.

k	1	2	4	8	16
CoM	0.51 (0.04)	0.55 (0.04)	0.75 (0.05)	0.89 (0.05)	0.97 (0.03)
KNN	0.50 (0.03)	0.52 (0.03)	0.61 (0.03)	0.75 (0.03)	0.90 (0.02)
KNN'	.	0.55 (0.03)	0.68 (0.03)	0.80 (0.02)	0.90 (0.02)
Kernel	0.54 (0.03)	0.52 (0.03)	0.64 (0.03)	0.77 (0.03)	0.92 (0.02)
Kernel'	.	0.58 (0.03)	0.69 (0.02)	0.81 (0.03)	0.92 (0.02)
GLM	0.36 (0.04)	0.41 (0.04)	0.49 (0.04)	0.57 (0.03)	0.65 (0.03)
GLM'	.	0.35 (0.04)	0.36 (0.04)	0.36 (0.05)	0.35 (0.05)
Tree	0.44 (0.03)	0.44 (0.03)	0.45 (0.03)	0.50 (0.03)	0.55 (0.03)
Tree'	.	0.46 (0.03)	0.51 (0.04)	0.51 (0.07)	0.47 (0.07)
CoM	0.46 (0.05)	0.50 (0.06)	0.63 (0.05)	0.75 (0.08)	0.85 (0.07)
KNN	0.46 (0.03)	0.49 (0.03)	0.57 (0.03)	0.67 (0.03)	0.78 (0.02)
KNN'	.	0.50 (0.03)	0.64 (0.03)	0.77 (0.02)	0.87 (0.02)
Kernel	0.50 (0.03)	0.48 (0.03)	0.57 (0.03)	0.68 (0.03)	0.80 (0.02)
Kernel'	.	0.53 (0.03)	0.66 (0.03)	0.77 (0.02)	0.85 (0.02)
GLM	0.35 (0.03)	0.41 (0.04)	0.46 (0.04)	0.53 (0.03)	0.58 (0.04)
GLM'	.	0.35 (0.03)	0.36 (0.04)	0.36 (0.04)	0.36 (0.06)
Tree	0.42 (0.03)	0.44 (0.03)	0.45 (0.03)	0.46 (0.03)	0.49 (0.03)
Tree'	.	0.43 (0.03)	0.47 (0.04)	0.48 (0.06)	0.46 (0.08)

CoM, concentration of measure; KNN, k-nearest-neighbours; GLM, generalized linear model; Tree, regression tree.

Table 3: Percentage of correct classification of the five phonemes against the five methods. The highest percentage of each column is marked in bold.

	/a/	/ɔ/	/d/	/i/	/ʃ/
CoM	0.769	0.768	0.966	0.985	0.994
KNN	0.724	0.791	0.985	0.974	1.000
Kernel	0.720	0.805	0.984	0.972	0.999
GLM	0.790	0.723	0.982	0.959	0.992
Tree	0.708	0.694	0.956	0.878	0.926

training curves. The results are detailed in Table 3. Our concentration of measure approach only uniformly outperforms the regression tree classifier, but has comparable performance to the other three methods, and none of the competing methods uniformly outperforms ours.

5.4 The expectation-maximization algorithm in practice

The experiments described and depicted below make use of the trace norm only. It was determined through experimentation that the expectation-maximization algorithm we propose in Section 4.3 does not perform well under the topology of either the Hilbert-Schmidt or operator norms as they give more emphasis to the principle eigenvalue at the expense of the others. The usual behavior under these norms is for all estimates to converge to the average of all of the data points. This is in contrast to the better performance of the algorithm making use of the trace norm, which is somewhat more uniform in its treatment of the eigenstructure.

As a first test case, this algorithm was run given three target covariance operators, which were constructed by taking three randomly generated orthonormal bases U_i and a diagonal operator D of eigenvalues decaying at a rate $\lambda_k = O(k^{-4})$ and multiplying $\Sigma_i = U_i D U_i^T$. Let the three target covariance operators be denoted as Σ_a , Σ_b , and Σ_c . For each of these operators, 1000 rank four data points were generated from a zero mean Gaussian process. From the data, the algorithm initializes three estimates $\hat{\Sigma}_1^{(t)}$, $\hat{\Sigma}_2^{(t)}$, and $\hat{\Sigma}_3^{(t)}$, which attempt to locate the three target operators as the method iterates.

Table 4: Clustering of simulated operators.

	Rank 4 Operators			Rank 1 Operators		
	Cluster 1	Cluster 2	Cluster 3	Cluster 1	Cluster 2	Cluster 3
Label a	0	0	1000	169	557	274
Label b	0	1000	0	184	549	267
Label c	1000	0	0	483	230	287

Table 5: Clustering 2000 phoneme curves into 5 clusters.

	Cluster 1	Cluster 2	Cluster 3	Cluster 4	Cluster 5
/a/	3	0	397	0	0
/ɔ/	0	0	397	1	2
/d/	0	0	0	378	22
/i/	212	0	0	1	187
/f/	4	393	0	0	3

After 15 iterations, the original 3000 data points were perfectly separated into three groups. To make the problem harder, a second test case was run identical to the first except that the observed operators are all of rank one. The resulting clusters from both tests are in Table 4. The inaccuracy in the rank one setting is equivalent to the poor performance of classification of rank one operators detailed in Tables 1 and 2.

For the phoneme data all 400 sample curves from each of the five phoneme sets were clustered individually as curves. The algorithm was run for 20 iterations and told to partition the data into five clusters. The results are in Table 5. Cluster 3 captured almost all of the vowels /a/ and /ɔ/ , which, recalling their definition in Section 5.1, are quite similar in sound. Clusters 2 and 4 contain the majority of /f/ and /d/ curves, respectively. Lastly, Clusters 1 and 5 split the set of /i/ curves roughly in half. Rerunning the experiment with only four clustered nicely separated the data as displayed in Table 6. Cluster 3 again contains almost all of the /a/ and /ɔ/ . The remaining clusters mostly partition the /d/ , /i/ , and /f/ with high accuracy. Hence, barring the ability to distinguish the very similar /a/ and /ɔ/ curves, the proposed expectation-maximization algorithm is an effective method for the unsupervised clustering of phonemes.

References

- Christophe Abraham, Pierre-André Cornillon, ERIC Matzner-Løber, and Nicolas Molinari. Unsupervised curve clustering using b-splines. *Scandinavian journal of statistics*, 30(3):581–595, 2003.
- Sylvain Arlot, Gilles Blanchard, and Etienne Roquain. Some nonasymptotic results on resampling in high dimension, i: confidence regions. *The Annals of Statistics*, 38(1):51–82, 2010.
- Peter L Bartlett and Shahar Mendelson. Rademacher and gaussian complexities: Risk bounds and structural results. *The Journal of Machine Learning Research*, 3:463–482, 2003.
- Peter L Bartlett, Stéphane Boucheron, and Gábor Lugosi. Model selection and error estimation. *Machine Learning*, 48 (1-3):85–113, 2002.
- Alain Berlinet, Gérard Biau, and Laurent Rouviere. Functional supervised classification with wavelets. In *Annales de l’ISUP*, volume 52, 2008.

Table 6: Clustering 2000 phoneme curves into four clusters.

	Cluster 1	Cluster 2	Cluster 3	Cluster 4
/a/	0	0	400	0
/ɔ/	1	0	398	1
/d/	384	1	0	15
/i/	1	5	0	394
/f/	0	397	0	3

- Denis Bosq. *Linear processes in function spaces: theory and applications*, volume 149. Springer Science & Business Media, 2012.
- Stéphane Boucheron, Gábor Lugosi, and Pascal Massart. *Concentration inequalities: A nonasymptotic theory of independence*. Oxford University Press, 2013.
- George Casella and Roger L Berger. *Statistical inference*, volume 2. Duxbury Pacific Grove, CA, 2002.
- Chung Chang, Yakuan Chen, and R Todd Ogden. Functional data classification: a wavelet approach. *Computational Statistics*, 29(6):1497–1513, 2014.
- Aurore Delaigle and Peter Hall. Achieving near perfect classification for functional data. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)*, 74(2):267–286, 2012.
- Zhou Fan. Confidence regions for infinite-dimensional statistical parameters. *Part III essay in Mathematics, University of Cambridge*, 2011. <http://web.stanford.edu/~zhoufan/PartIIIEssay.pdf>.
- Frédéric Ferraty and Philippe Vieu. Curves discrimination: a nonparametric functional approach. *Computational Statistics & Data Analysis*, 44(1):161–173, 2003.
- Frédéric Ferraty and Philippe Vieu. *Nonparametric functional data analysis: theory and practice*. Springer Science & Business Media, 2006.
- Stefan Fremdt, Josef G Steinebach, Lajos Horváth, and Piotr Kokoszka. Testing the equality of covariance operators in functional samples. *Scandinavian Journal of Statistics*, 40(1):138–152, 2013.
- Evarist Giné and Richard Nickl. Confidence bands in density estimation. *The Annals of Statistics*, 38(2):1122–1170, 2010.
- Evarist Giné and Richard Nickl. *Mathematical Foundations of Infinite-Dimensional Statistical Models*. Cambridge University Press, 2016.
- Richard H Glendinning and RA Herbert. Shape classification using smooth principal components. *Pattern recognition letters*, 24(12):2021–2030, 2003.
- Peter Hall, DS Poskitt, and Brett Presnell. A functional dataanalytic approach to signal discrimination. *Technometrics*, 43(1):1–9, 2001.
- Trevor Hastie, Andreas Buja, and Robert Tibshirani. Penalized discriminant analysis. *The Annals of Statistics*, pages 73–102, 1995.
- Lajos Horváth and Piotr Kokoszka. *Inference for functional data with applications*, volume 200. Springer Science & Business Media, 2012.
- Leon Isserlis. On a formula for the product-moment coefficient of any order of a normal frequency distribution in any number of variables. *Biometrika*, pages 134–139, 1918.
- Gareth M James and Trevor J Hastie. Functional linear discriminant analysis for irregularly sampled curves. *Journal of the Royal Statistical Society. Series B, Statistical Methodology*, pages 533–550, 2001.
- Ci-Ren Jiang, John AD Aston, and Jane-Ling Wang. A functional approach to deconvolve dynamic neuroimaging data. *arXiv preprint arXiv:1411.2051*, 2014.
- Gerard Kerkycharian, Richard Nickl, and Dominique Picard. Concentration inequalities and confidence bands for needlet density estimators on compact homogeneous manifolds. *Probability Theory and Related Fields*, 153(1-2):363–404, 2012.
- Vladimir Koltchinskii. Rademacher penalties and structural risk minimization. *Information Theory, IEEE Transactions on*, 47(5):1902–1914, 2001.
- Vladimir Koltchinskii. Local rademacher complexities and oracle inequalities in risk minimization. *The Annals of Statistics*, 34(6):2593–2656, 2006.
- Michel Ledoux. *The concentration of measure phenomenon*, volume 89. American Mathematical Soc., 2001.

- Karim Lounici and Richard Nickl. Global uniform risk bounds for wavelet deconvolution estimators. *The Annals of Statistics*, 39(1):201–231, 2011.
- Hans-Georg Müller and Ulrich Stadtmüller. Generalized functional linear models. *Annals of Statistics*, pages 774–805, 2005.
- Victor M Panaretos, David Kraus, and John H Maddocks. Second-order comparison of gaussian random functions and the geometry of dna minicircles. *Journal of the American Statistical Association*, 105(490):670–682, 2010.
- Jie Peng and Hans-Georg Müller. Distance-based clustering of sparsely observed stochastic processes, with applications to online auctions. *The Annals of Applied Statistics*, pages 1056–1077, 2008.
- Davide Pigoli, John AD Aston, Ian L Dryden, and Piercesare Secchi. Distances and inference for covariance operators. *Biometrika*, page asu008, 2014.
- Davide Pigoli, Pantelis Z Hadjipantelis, John S Coleman, and John AD Aston. The analysis of acoustic phonetic data: exploring differences in the spoken romance languages. *arXiv preprint arXiv:1507.07587*, 2015.
- James O Ramsay and Bernard W Silverman. *Functional data analysis*. New York: Springer, 2005.
- Walter Rudin. *Real and complex analysis*. Tata McGraw-Hill Education, 1987.
- Michel Talagrand. New concentration inequalities in product spaces. *Inventiones mathematicae*, 126(3):505–563, 1996.