

FUNCTIONAL DATA CLASSIFICATION AND CLUSTERING
RESEARCH ARTICLE

A classification method for functional data based on extremality indices

CATALINA LESMES^{1,2,3,*} , FRANCISCO ZULUAGA¹ , HENRY LANIADO⁴ ,
ANDRES GOMEZ¹ , and ANDREA CARVAJAL¹ 

¹School of Applied Sciences and Engineering, Universidad EAFIT, Medellín, Colombia

²Institute of Experimental and Applied Physics, Czech Technical University, Prague, Czech Republic

³Faculty of Nuclear Sciences and Physical Engineering, Czech Technical University, Prague, Czech Republic

⁴Department of Mathematics, Universidad del Cauca, Popayán, Colombia

(Received: 23 December 2025 · Accepted in final form: 03 February 2026)

Abstract

Functional data arise when each observation is a curve or trajectory. Despite the progress of classifiers based on depth, practitioners often need simple and interpretable tools that work well when group differences are driven by envelope extremality. This article introduces a new representation based on extremality for supervised classification: the extremality–extremality plot, which maps each curve to two class-referenced indices that quantify, for each class, the average proportion of time that the class curves lie below or above the given curve. Our objectives are to (i) motivate and formalize this representation and its relationship to approaches based on depth, (ii) describe a practical training and evaluation protocol that avoids optimistic bias, and (iii) assess performance on controlled synthetic scenarios and on three benchmark real datasets (Berkeley growth study, Tecator spectrometric dataset, and mitochondrial calcium overload). Methods include linear and quadratic discriminant analysis, k-nearest neighbors, support vector machines, and random forests trained on the two-dimensional extremality features. Across experiments, the proposed classifier attains high accuracy when classes are well separated and remains competitive under overlap, often with stable performance across resampling folds. These results indicate that projections based on extremality provide a simple, interpretable, and, in practice, computationally light alternative for supervised classification of functional data.

Keywords: epigraph · extremality · functional data · hypograph · supervised classification

Mathematics Subject Classification: Primary 62R10 · Secondary 62H30

1. INTRODUCTION

Data processing and analysis have become essential for understanding complex phenomena across science, industry, and the public sector. As measurements are increasingly recorded as high-frequency trajectories, it is natural to treat each observation as a function. Functional data analysis (FDA) provides the statistical framework for modeling, inference, and decision making with such curve-valued observations.

*Corresponding author. Email: clesmes@eafit.edu.co (C. Lesmes)

Early formalizations of this paradigm go back to [1]. Since then, a broad methodology has been developed to address the specific features of functional data. Classical summaries such as the mean and the covariance have been adapted to sparse and irregular designs [2]. More advanced procedures for classification [3], clustering [4], dimension reduction and inference [5], and two-sample assessment via homogeneity tests [6] have also been extended to the functional setting. Recent contributions further illustrate the breadth of applications, including functional geostatistics based on Andrews curves [7] and robust nonparametric regression with functional covariates [8]. Related advances in high-dimensional model selection can improve inference in complex functional contexts [9].

This article focuses on supervised classification. Depth-based approaches have proven effective for functional data [10], most notably the depth–depth classifier (DD) [11] and its multiclass extension DD^G [12]. We adopt the same embedding idea, mapping curves to a low-dimensional, interpretable feature space, and replace statistical depth with extremality.

Extremality for functional data [13] provides a natural top-to-bottom ordering of curves via the hypograph and epigraph indices. These indices underpin several graphical and inferential tools, including the outliergram [14], a functional boxplot variant based on epigraphs and hypographs [15] (the classic functional boxplot is based on band depth [16]), rank-based two-sample tests [17], and efficient clustering procedures [18].

Our contribution is twofold. First, we introduce the extremality–extremality (EE) plot, which maps each curve to two class-referenced extremality indices and thus to a two-dimensional space. Second, we define a companion EE classifier that applies standard decision rules in this space, including linear and quadratic discriminant analysis, k-nearest neighbors (kNN), support vector machines (SVM), and random forests (RF). We compare EE with DD^G on synthetic scenarios and on three benchmark real datasets: the Berkeley growth study [19], the Tecator spectrometric dataset [20], and the mitochondrial calcium overload (MCO) dataset [21], using an evaluation protocol that recomputes extremality features for validation data with respect to the corresponding training folds to avoid optimistic bias.

The remainder of the article is organized as follows. Section 2 reviews the DD and DD^G frameworks and the epigraph and hypograph indices. In Section 3, we introduce the EE plot and the EE classifier. In Section 4, empirical results on simulated and real datasets are reported. Section 5 provides our conclusions. Additional results are provided in Appendix A.

2. BACKGROUND

We briefly review the concepts underlying our approach. The proposed classifier builds on the DD framework of [10, 11] and on extremality notions based on epigraphs and hypographs. We first recall DD classification and several functional depths commonly used to construct DD plots, and then summarize the epigraph and hypograph indices.

2.1 Depth–depth classifier and its generalization

The DD classifier [11] and its multiclass extension DD^G [12] rely on the DD plot [22], which maps each observation to the plane by pairing its depths with respect to two reference distributions (or two class samples). This produces a low-dimensional, interpretable representation in which standard decision rules can be applied.

For the DD plot, one may use a variety of depth notions. We recall the following three popular choices for functional data:

- Fraiman–Muniz (FM) depth. Given a sample of functions $x_1, \dots, x_N$ observed on an interval $\mathcal{I} = [0, T]$, let $S_t = (x_1(t), \dots, x_N(t))$ be the vector of values at time t , let $F_{N,t}$ denote its empirical cumulative distribution function, and let $Z_i(t)$ be a univariate depth of $x_i(t)$ relative to S_t . The FM depth of the i th curve is the time average stated as

$$\text{FM}_i = \frac{1}{T} \int_0^T Z_i(t) dt.$$

In general, $\text{FM}_i \in [0, 1]$ whenever $Z_i(t) \in [0, 1]$. Under the choice $Z_i(t) = 1 - |1/2 - F_{N,t}(x_i(t))|$ [23], the range tightens to $\text{FM}_i \in [1/2, 1]$. Other univariate depth notions (for example, univariate halfspace/Tukey depth) can be used in place of $1 - |1/2 - F_{N,t}(x)|$.

- h-mode (hM) depth. Following [24], the population h -mode depth of a datum x_0 is defined by

$$f_h(x_0) = \mathbb{E} \left(K \left(\frac{m(x_0, X)}{h} \right) \right),$$

where X denotes a random curve from the population, m is a suitable metric (or semi-metric), K is a nonnegative kernel typically nonincreasing in m/h , and $h > 0$ is a bandwidth. For a sample $X_1, \dots, X_N$ with observed values $x_1, \dots, x_N$, the empirical version is expressed as

$$\hat{f}_h(x_0) = \frac{1}{N} \sum_{i=1}^N K \left(\frac{m(x_0, x_i)}{h} \right).$$

- Random projection (RP) depth. Also in [24], one projects the data onto random directions and averages a univariate depth across projections. If $a_1, \dots, a_R$ are random directions and $D_{a_r}(x)$ denotes the univariate depth of the projection of x along a_r , the RP depth is given by

$$\text{RP}(x) = \frac{1}{R} \sum_{r=1}^R D_{a_r}(x),$$

with $R = 50$ a common default. In multivariate or functional settings, one may also aggregate componentwise depths via suitable weights.

Figure 1 illustrates typical DD plots in two bivariate normal scenarios: identical distributions (panel a) concentrate near the diagonal, while mean-shifted distributions (panel b) depart from it.

The formal definition of the DD plot is presented in [22]. Let F and G be two probability distributions on $\mathbb{R}^d$, and let D denote a depth functional. The DD plot associated with (F, G) is the set of pairs of depths of the same point with respect to F and G established as

$$\text{DD}(F, G) = \{ (D_F(x), D_G(x)), x \in \mathbb{R}^d \}.$$

When F is unknown and a sample $X_1, \dots, X_n$ is observed, let F_n be the empirical distribution of $X_1, \dots, X_n$. To assess agreement with a specified distribution G , one considers the empirical DD plot stated as

$$\text{DD}(F_n, G) = \{ (D_{F_n}(X_i), D_G(X_i)), i = 1, \dots, n \}.$$

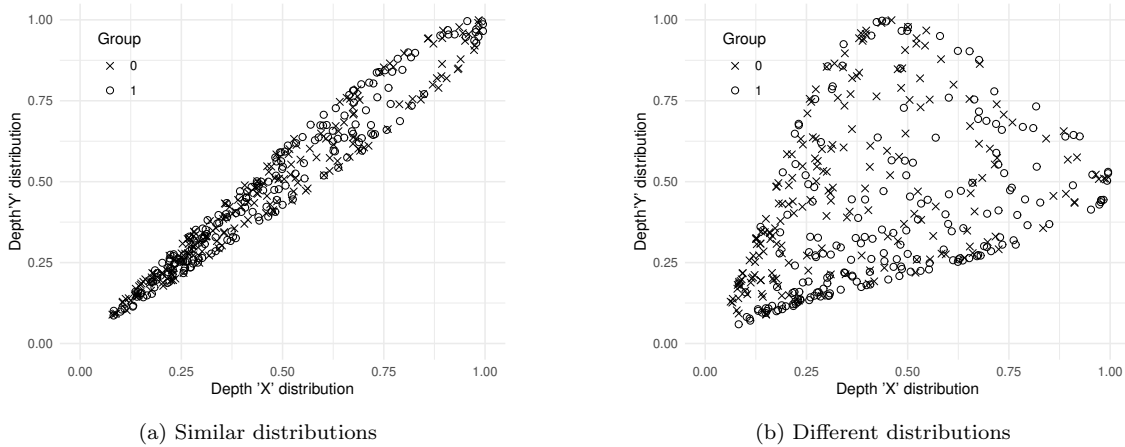


Figure 1. DD plots for two samples $\mathcal{X}$ and $\mathcal{Y}$, where, when $\mathcal{X}$ and $\mathcal{Y}$ come from the same distribution, points cluster along the diagonal.

In the two-sample setting, with samples $\mathcal{X} = \{X_1, \dots, X_n\}$ from F and $\mathcal{Y} = \{Y_1, \dots, Y_m\}$ from G , let F_n and G_m be their empirical distributions. Then, the pooled DD plot is given by

$$\text{DD}(F_n, G_m) = \{(D_{F_n}(z), D_{G_m}(z)), z \in \mathcal{X} \cup \mathcal{Y}\}. \quad (1)$$

Figure 1 illustrates typical patterns using bivariate normal data generated in the R software [25] by the function `mvrnorm` from the MASS package [26]. Panel (a) compares two identical distributions with identity covariance, while panel (b) compares two distributions with the same covariance and different means. In the first case, points concentrate near the diagonal of the DD plot. In the second case, they deviate from it.

As documented in [22], departures from the diagonal convey distributional differences: a location shift often produces a heart-shaped pattern, a scale difference tends to yield an arch above the diagonal, skewness may induce an asymmetric arch, and kurtosis differences can shift the lower portion of the cloud to one side. These visual cues make the DD plot a compact diagnostic for comparing distributions.

This representation also enables classification. The DD classifier maps curves to the DD plane and selects a separating rule in that plane so as to minimize misclassification error on the training data [11]. The multiclass extension DD^G embeds each observation into a vector of depths with respect to all classes. Then, it applies off-the-shelf classifiers such as linear or quadratic discriminant analysis, kNN, or SVM [12]. Unlike DD, which is binary by construction, DD^G naturally handles more than two groups.

2.2 Epigraph and hypograph indices

Extremality for functions is based on the notions of hypograph and epigraph. Let $C(\mathcal{I})$ be the space of real-valued continuous functions on a compact interval $\mathcal{I} \subset \mathbb{R}$. For $x \in C(\mathcal{I})$, write its graph as $G(x) = \{(t, x(t)), t \in \mathcal{I}\}$. Following [27], the hypograph and epigraph of x are the sets stated as $\text{hyp}(x) = \{(t, y) \in \mathcal{I} \times \mathbb{R}, y \leq x(t)\}$ and $\text{epi}(x) = \{(t, y) \in \mathcal{I} \times \mathbb{R}, y \geq x(t)\}$. Epigraph and hypograph indices were introduced in [27] to induce an ordering of curves by extremality. Given a sample $X_1, \dots, X_n \in C(\mathcal{I})$ with observations $x_1, \dots, x_n$, and writing $I(A)$ for the indicator of event A , the hypograph index (HI) and the epigraph index (EI) of a reference curve x are defined

$$\text{HI}(x) = \frac{1}{n} \sum_{i=1}^n I(G(x_i) \subset \text{hyp}(x)) = \frac{1}{n} \sum_{i=1}^n I(x_i(t) \leq x(t), t \in \mathcal{I})$$

and

$$\text{EI}(x) = \frac{1}{n} \sum_{i=1}^n I(G(x_i) \subset \text{epi}(x)) = \frac{1}{n} \sum_{i=1}^n I(x_i(t) \geq x(t), t \in \mathcal{I}).$$

Because HI and EI require containment over the whole interval, they may lose resolution when curves are very close or widely spaced [14]. To address this, modified versions were proposed in [14, 15]. For a reference sample $\mathcal{S} = \{x_1, \dots, x_n\}$, the modified hypograph index (MHI) and modified epigraph index (MEI) are given by

$$\begin{aligned} \text{MHI}_{\mathcal{S}}(x) &= \frac{1}{n \lambda(\mathcal{I})} \sum_{i=1}^n \lambda\{t \in \mathcal{I}, x_i(t) \leq x(t)\}, \\ \text{MEI}_{\mathcal{S}}(x) &= \frac{1}{n \lambda(\mathcal{I})} \sum_{i=1}^n \lambda\{t \in \mathcal{I}, x_i(t) \geq x(t)\}, \end{aligned} \tag{2}$$

where λ denotes Lebesgue measure on $\mathbb{R}$. A basic relationship, used throughout this article, is that for $x \notin \mathcal{S}$ (ties of null measure), one has $\text{MHI}_{\mathcal{S}}(x) + \text{MEI}_{\mathcal{S}}(x) = 1$, whereas for $x = x_k \in \mathcal{S}$ under non-strict inequalities, the sum equals $1 + 1/n$ [14, 15]. In practice, when curves are observed on a discrete grid, the integrals stated in (2) are approximated by Riemann sums. We adopt non-strict comparisons (" $\leq$ " for MHI and " $\geq$ " for MEI), counting exact ties according to those inequalities. This deterministic convention preserves the out-of-sample relation $\text{MHI} + \text{MEI} = 1$ up to discretization error, and yields $1 + 1/n$ (or $1 + 1/m$) in-sample on the corresponding axis.

3. EXTREMALITY-EXTREMALITY PLOT AND CLASSIFIER

As indicated in Section 2, the MHI and MEI order curves and remain informative even when trajectories are closely spaced. These indices have supported graphical and inferential tools for comparison and clustering of functional samples [17, 18]. In this section, we introduce a visualization tailored to supervised problems: the EE plot, together with a companion classifier that operates in the resulting two-dimensional space.

3.1 Extremality-extremality plot

The EE plot displays, on each axis, an extremality index of a curve computed with respect to one of the two class samples. Let $\mathcal{X} = \{X_1, \dots, X_n\}$ and $\mathcal{Y} = \{Y_1, \dots, Y_m\}$ be samples in $C(\mathcal{I})$ from two classes. By analogy with the DD construction given in (1), we define EE plots using either MHI on both axes or MEI on both axes as

$$\begin{aligned} \text{EE}_H(\mathcal{X}, \mathcal{Y}) &= \{(\text{MHI}_{\mathcal{X}}(z), \text{MHI}_{\mathcal{Y}}(z)), z \in C(\mathcal{I})\}, \\ \text{EE}_E(\mathcal{X}, \mathcal{Y}) &= \{(\text{MEI}_{\mathcal{X}}(z), \text{MEI}_{\mathcal{Y}}(z)), z \in C(\mathcal{I})\}. \end{aligned}$$

In practice, we display the pooled EE plot by evaluating these mappings only on the observed curves as

$$\begin{aligned} \text{EE}_H^{\text{pool}}(\mathcal{X}, \mathcal{Y}) &= \{(\text{MHI}_{\mathcal{X}}(z), \text{MHI}_{\mathcal{Y}}(z)), z \in \mathcal{X} \cup \mathcal{Y}\}, \\ \text{EE}_E^{\text{pool}}(\mathcal{X}, \mathcal{Y}) &= \{(\text{MEI}_{\mathcal{X}}(z), \text{MEI}_{\mathcal{Y}}(z)), z \in \mathcal{X} \cup \mathcal{Y}\}. \end{aligned}$$

Each axis is computed with respect to one reference sample (either $\mathcal{X}$ or $\mathcal{Y}$). Therefore, as noted in Section 2, for any point z that does not belong to the sample used on a given axis (ties have probability zero under continuity), we have $\text{MEI} = 1 - \text{MHI}$ on that axis. If z belongs to that sample and non-strict inequalities are used, the relation becomes $\text{MEI} = 1 - \text{MHI} + 1/n$ (or $+1/m$ on the $\mathcal{Y}$ axis). Consequently, EE_H and EE_E are related by a central symmetry: away from in-sample points, the mapping in the EE plane is $(u, v) \mapsto (1 - u, 1 - v)$, that is, a 180° rotation about $(1/2, 1/2)$. At in-sample points, small axis-specific shifts of size $1/n$ or $1/m$ apply.

The qualitative geometry of EE plots mirrors distributional differences between classes. When the two samples follow similar distributions, points concentrate along the 45° line. For sinusoidal signals, a chromosome-like pattern may appear. As amplitude or dispersion increases, the cloud spreads. When classes hardly overlap, L-shaped clusters emerge near the corners of the unit square. With mild overlap, these corners become rounded. These behaviors are illustrated in Figures 2 and 3, which display increasing separation along the axes and pronounced corner structures when classes are well separated.

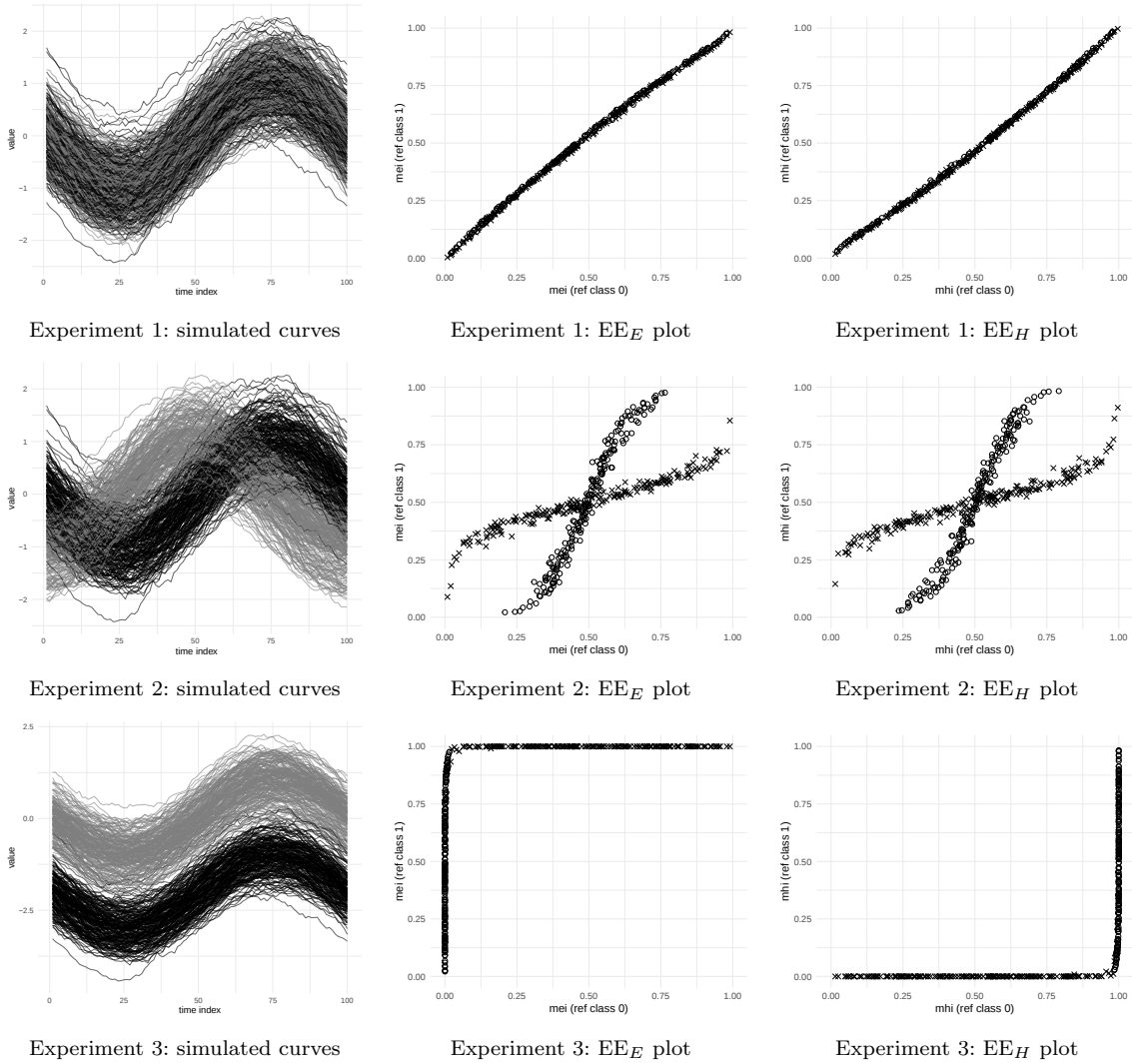


Figure 2. Experiments 1–3: (left) simulated curves, with class 0 in black and class 1 in gray; (center) EE_E plot; and (right) EE_H plot with markers in EE plots given by class 0 = $\times$, class 1 = $\circ$, and both axes are scaled to $[0, 1]$.

3.2 Extremality-extremality classifier

The EE plot embeds each curve into $\mathbb{R}^2$ by pairing its extremality indices with respect to the two class samples. In this plane, standard decision rules can separate groups, in the same spirit as DD methods [11, 12]. We consider linear (LDA) and quadratic (QDA) discriminant analyses, kNN, SVM, and RF.

The EE classifier proceeds as follows. Given two classes $\mathcal{X}$ and $\mathcal{Y}$ and a chosen index (MHI or MEI), each curve z is mapped to the pair of indices computed with respect to $\mathcal{X}$ and to $\mathcal{Y}$, yielding a two-dimensional feature vector. A classifier is then trained on these features. To avoid optimistic bias, when cross-validating, we recompute the indices for validation curves with respect to the corresponding training folds only. Predictions for held-out curves use the same training references. Accuracy is the primary performance metric reported in Section 4.

Because MEI and MHI are tightly related (outside the reference sample, $\text{MEI} = 1 - \text{MHI}$ on the same axis, with a small in-sample shift), EE plots based on MEI and on MHI convey nearly the same geometric information.

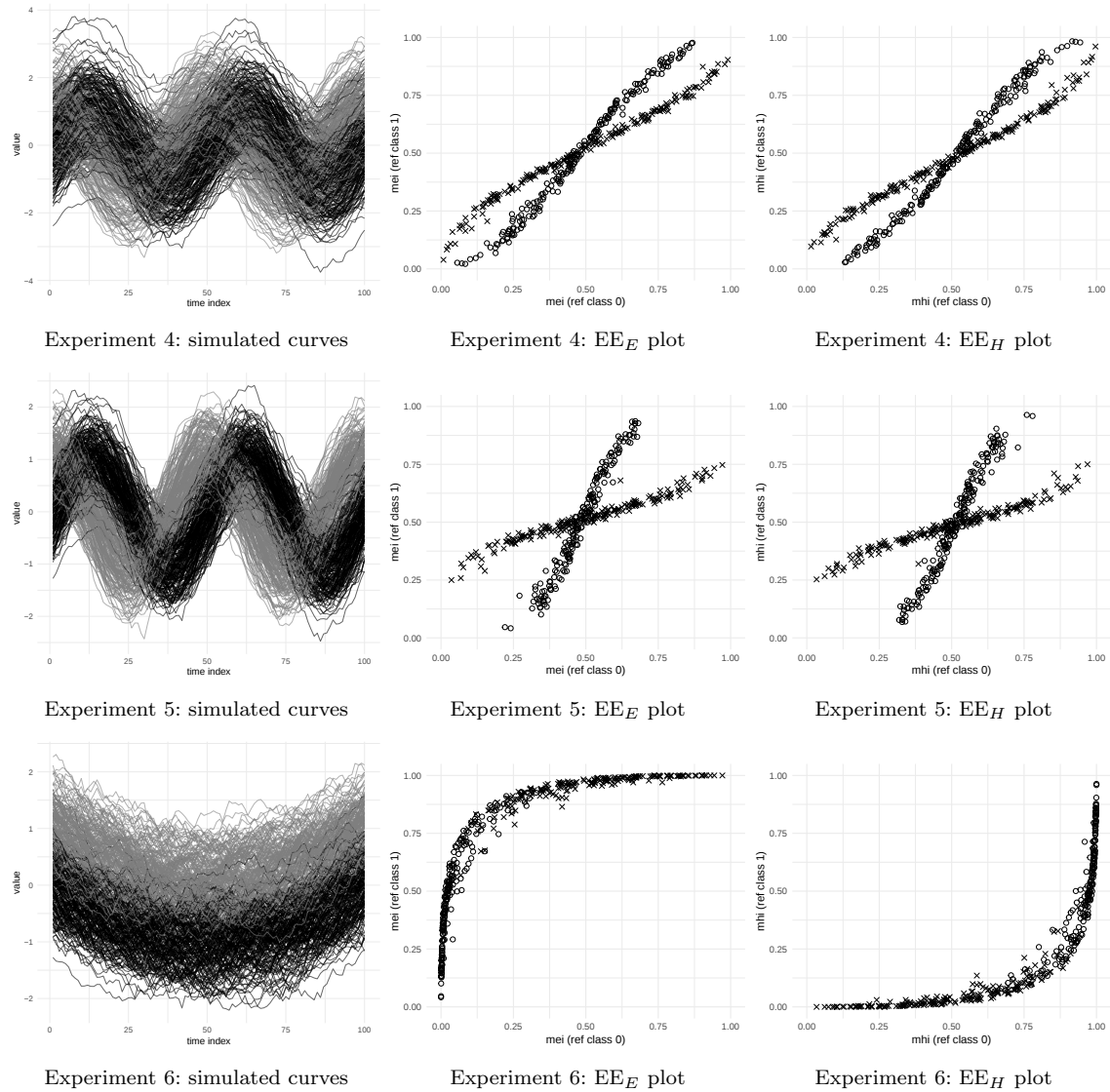


Figure 3. Experiments 4–6: (left) simulated curves with class 0 in black and class 1 in gray; (center) EE_E plot; and (right) EE_H plot with markers in EE plots class 0 = x, class 1 = o, and both axes are scaled to $[0, 1]$.

4. APPLICATIONS

We evaluate the EE classifier on both synthetic and real datasets. For each study, we report accuracy under cross-validation (CV) and summarize its distribution across folds using the median and the interquartile range (IQR). To avoid optimistic bias, extremality features for validation/test curves are always computed with respect to the corresponding training folds only.

4.1 Synthetic data

The synthetic scenarios correspond to the distributions depicted in Figures 2 and 3. Data were generated with the function `generate_gauss_fdata` from the `roahd` package [28] (version 1.4.3) using a fixed random seed (`set.seed(45)`). Each curve is observed on an equally spaced grid of 100 points over $t \in [-1, 1]$ and is a realization of a Gaussian process with mean function $\mu(t)$ and exponential covariance kernel $C(s, t) = \alpha \exp(-\beta |s - t|)$, where $\alpha > 0$ controls amplitude and $\beta > 0$ controls the rate of decay.

The function `generate_gauss_fdata` produces univariate functional datasets from user-specified mean functions and a covariance operator (the latter specified via `exp_cov_function`). We designed six experiments to probe the behavior of the EE classifier under different characteristics: identical distributions (Experiment 1), complementary mean functions such as sine and cosine (Experiment 2), non-overlapping classes (Experiment 3), larger amplitude (Experiment 4), shorter correlation length (higher roughness, Experiment 5), and distinct polynomial centerlines (Experiment 6). The mean functions and covariance parameters are summarized in Table 1. In all experiments, each class contains $n = 200$ curves with 100 time points per curve.

We trained classifiers using an outer stratified K -fold CV at the curve level ($K = 100$). In each outer fold, MEI/MHI and depth features for validation curves were recomputed with respect to the corresponding training classes only. Models were then fitted on the training features and evaluated on the held-out folds. Hyperparameters were tuned by inner 5-fold CV using the default grids of the `caret` package [29], with `set.seed(7)` for reproducibility. For each method, we report the distribution of outer-fold accuracies (minimum, quartiles, median, mean, and maximum), and we emphasize medians and IQRs as primary summaries.

Table 1. Synthetic scenarios: mean functions and covariance parameters used for data generation.

Experiment	Mean functions $\mu_1(t), \mu_2(t)$	α	β
1	$\mu_1(t) = \sin(\pi t), \mu_2(t) = \sin(\pi t)$	0.2	0.3
2	$\mu_1(t) = \sin(\pi t), \mu_2(t) = \cos(\pi t)$	0.2	0.3
3	$\mu_1(t) = \sin(\pi t) - 2, \mu_2(t) = \sin(\pi t)$	0.2	0.3
4	$\mu_1(t) = \sin(2\pi t), \mu_2(t) = \cos(2\pi t)$	0.7	0.3
5	$\mu_1(t) = \sin(2\pi t), \mu_2(t) = \cos(2\pi t)$	0.2	0.9
6	$\mu_1(t) = t^2 - 1, \mu_2(t) = t^2$	0.2	0.9

As summarized in Figure 4, Experiment 1 (identical class distributions) yields accuracies close to chance (medians near 0.50) for all methods, with no consistent advantage for any single classifier. In Experiment 2 (sine versus cosine), most methods are near-perfect (medians at 1.00), while LDA on EE features remains substantially lower, reflecting non-linearity in the EE plane. Experiment 3 (well-separated classes) is solved by all methods (medians at 1.00).

In Experiments 4 and 5 (increased amplitude and increased roughness, respectively), QDA, kNN, SVM, and RF on EE features achieve very high accuracies (medians at or near 1.00), whereas LDA on EE_E can be markedly weaker (for example, median 0.50 in Experiment 4).

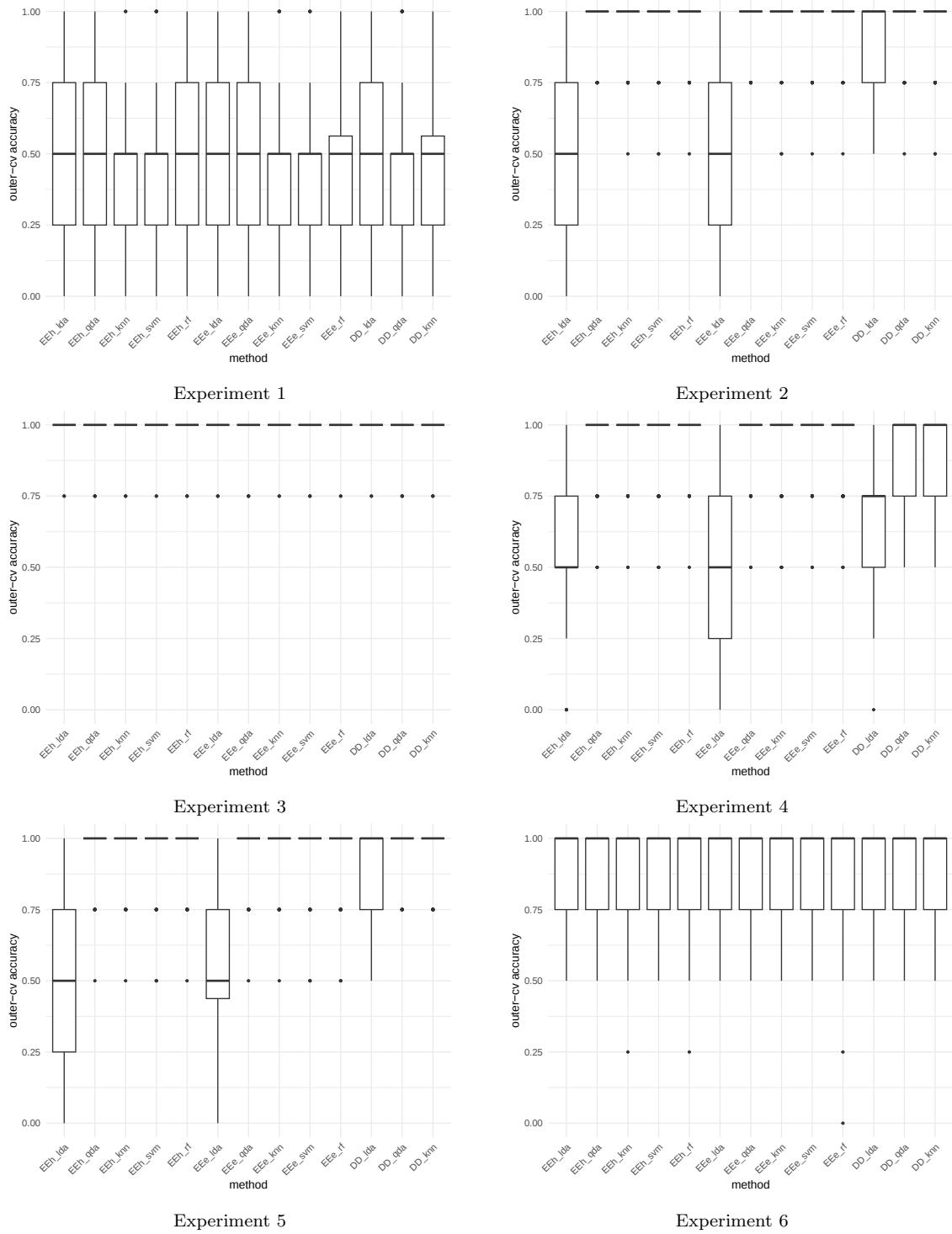


Figure 4. Boxplots of CV accuracy in synthetic experiments (stratified 100-fold outer CV, with validation features recomputed against training folds only) and set of classifiers.

Experiment 6 (distinct polynomial centerlines) shows uniformly high performance across methods (medians at 1.00 or close), with small differences in means and IQRs. Overall, EE_E and EE_H produce very similar patterns across all six scenarios, consistent with their near-equivalence.

4.2 Real data

We assess the EE classifier on three benchmark functional datasets available in the R package `fda.usc` [30] (version 2.2.0): Berkeley growth study [19], Tecator spectrometric dataset [20], and MCO dataset [21]. We fit models with stratified outer CV at the curve level (20-fold for Berkeley growth study, 50-fold for Tecator, and 30-fold for MCO). In every outer split, extremality and depth features for validation curves are recomputed using the training folds only. Inner 5-fold CV within the training folds is used for tuning, with `set.seed(7)` for reproducibility.

The Berkeley growth study [19] dataset comprises height trajectories for 39 boys and 54 girls from age 1 to 18 [19]. Figure 5 displays the curves by sex and the corresponding EE_E and EE_H plots. The two groups are broadly similar at early ages, with clearer separation at older ages where boys tend to be taller.

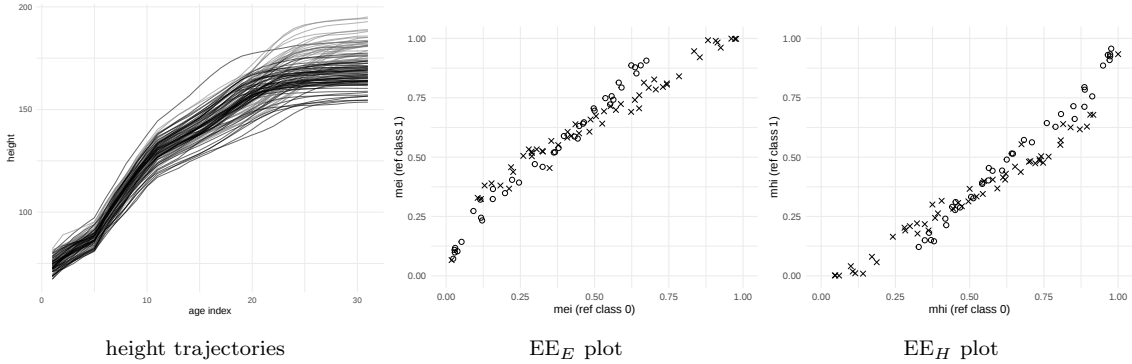


Figure 5. Curves and EE plots for Berkeley growth study data.

Figure 6 summarizes 20-fold outer CV accuracy across classifiers. Overall accuracies are around 75%, indicating moderate separability based on height trajectories. The EE_E features paired with LDA tend to be the weakest combination (mean accuracy around 0.54), whereas QDA with EE_H often yields one of the strongest summaries (median accuracy close to 0.80 and upper quartile at 1.00). The benchmark based on the FM depth with LDA/QDA/kNN is competitive, with DD+LDA also achieving median accuracy near 0.80.

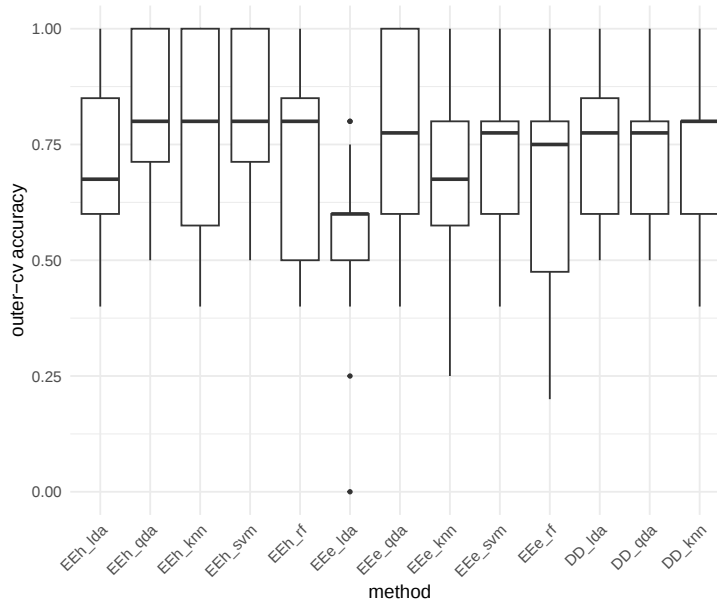


Figure 6. Boxplots of stratified 20-fold outer CV accuracy across classifiers (EE and DD) for Berkeley growth study data.

Remark. Because stratified 100-fold validation in the synthetic studies may leave only a handful of curves per validation fold, fold-wise accuracies of 100% can occur by chance. Therefore, we report medians and IQRs as more stable summaries. Future work could complement these results with repeated K-fold CV, such as 10×10 -fold CV, or bootstrap resampling to further stabilize the estimates.

The Tecator dataset [20] contains spectrometric absorbance curves from meat samples measured at 100 wavelengths between 850 and 1050 nm, grouped by fat content (low fat $< 20\%$ versus high fat $\geq 20\%$). Figure 7 shows the curves and the corresponding EE_E and EE_H plots. The absorbance curves for the two groups are strongly intertwined, resulting in substantial overlap in both the function space and the EE plane.

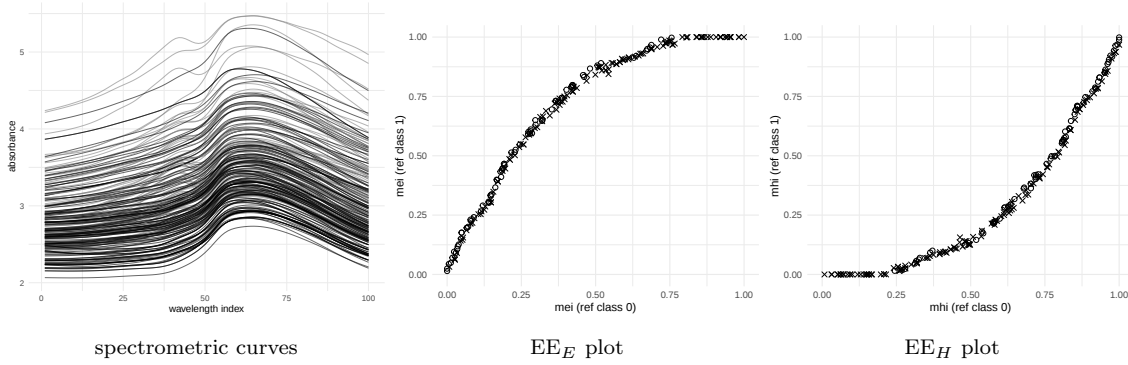


Figure 7. Curves and EE plots for Tecator data.

Figure 8 reports stratified 50-fold outer CV accuracy across classifiers. As expected under strong overlap, accuracies are moderate, with medians typically near 0.67 across most methods. The strongest median performance is often achieved by random forests on EE features (median ≈ 0.75), while LDA, QDA, and kNN on both EE and depth features cluster around 0.67, with similar IQRs. Overall, the EE classifier remains competitive and exhibits fold-to-fold stability comparable to that of the benchmark based on depth.

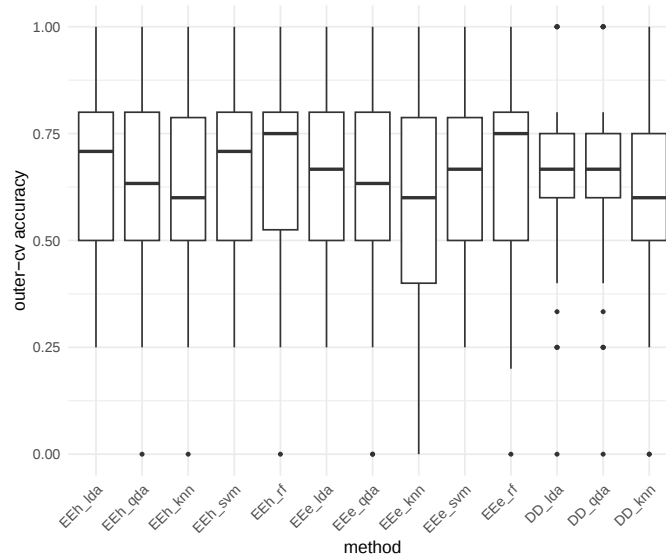


Figure 8. Boxplots of stratified 50-fold outer CV accuracy across classifiers (EE and DD) for Tecator data.

The MCO dataset records mitochondrial calcium overload in isolated mouse cardiac cells every 10 seconds over one hour, with two groups: control (no drug) and treatment (drug administered), where the drug is known to increase mitochondrial calcium loading [21]. Figure 9 shows the curves and their corresponding EE_E and EE_H representations. The two classes overlap markedly in function space, and the EE plots display only modest separation.

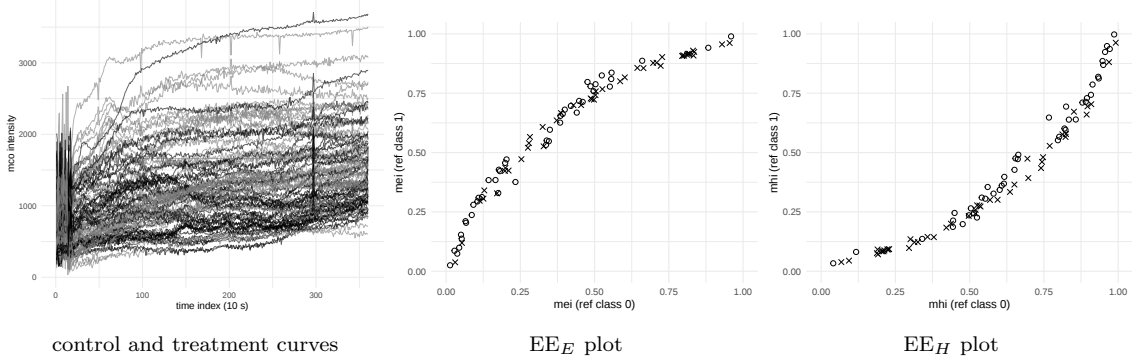


Figure 9. Curves and EE plots for the control and treatment groups with MCO data.

Figure 10 summarizes 30-fold outer CV accuracy across classifiers. Mean accuracies range from approximately 0.55 to 0.71, reflecting the limited separability of the groups. The highest mean accuracy is achieved by the benchmark based on depth, using kNN (mean ≈ 0.71). Among EE methods, QDA paired with either EE_E or EE_H attains similar median performance (median ≈ 0.67) with competitive fold-to-fold variability. Because some outer folds contain only a small number of validation curves, maxima at 1.00 can occur by chance. Therefore, medians and IQRs provide the most reliable indicators of performance.

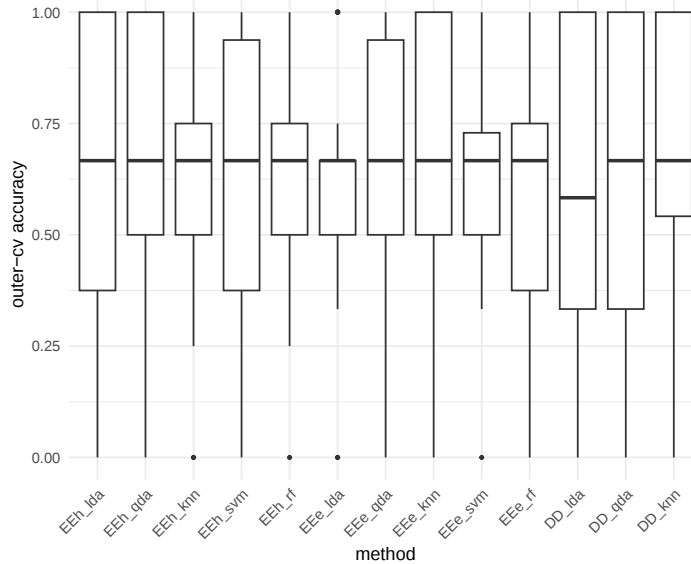


Figure 10. Boxplots of stratified 30-fold outer CV accuracy across classifiers (EE and DD) for MCO data.

5. CONCLUSIONS

We introduced a simple and interpretable representation for supervised classification of functional data: the EE plot, together with a companion EE classifier. The approach maps each curve to two class-referenced extremality features based on the modified epigraph and hypograph indices (MEI and MHI). Then, standard decision rules are applied in this two-dimensional space. We compared the EE classifier with a benchmark based on depth, namely a DD embedding (equivalently, DD^G in the binary case), under a common evaluation protocol. Across six synthetic scenarios, the empirical results follow intuition: when classes are identical, accuracies concentrate near chance. When classes are well separated, most methods approach perfect accuracy. Quadratic discriminant analysis frequently attains the best median performance, while linear discriminant analysis excels in settings that are close to linearly separable. The two EE variants, EE_E and EE_H , yield nearly indistinguishable outcomes, consistent with the near-identity between MEI and MHI on out-of-sample points and the small in-sample shifts discussed earlier. Fold-wise variability is often modest when classes are well separated but can increase under very small validation folds. This reinforces the use of medians and IQRs as primary summaries.

On real datasets, accuracies reflect the intrinsic difficulty of each task. For the Berkeley growth study, performance is moderate (approximately 0.55–0.83 across methods, with the best medians near 0.80), with EE often exhibiting comparable or tighter IQRs than alternatives. For Tecator, where groups strongly overlap, accuracies lie in the mid-to-high 0.60s (about 0.63–0.75). For the mitochondrial calcium overload data, accuracies are lower on average (mid-0.60s overall, with methods ranging roughly from 0.55 to 0.71), but the EE classifier remains competitive with DD embeddings based on depth and in some settings shows smaller dispersion across folds. Overall, the evidence supports EE as a practical alternative to projections based on depth.

This work has several natural extensions. First, multiclass problems can be addressed by constructing vectors of extremality features with respect to all classes, in direct analogy with DD^G . Second, larger datasets would allow a more granular assessment of stability and sensitivity to hyperparameters. Third, a systematic study of computational cost and the impact of discretization and ties would further clarify operating regimes where EE is most advantageous.

In summary, projections based on extremality provide a compact, low-dimensional view of functional samples that is easy to compute and interpret. The EE classifier delivers accuracy competitive with DD^G and exhibits stable performance across a range of synthetic and real scenarios, reinforcing its value as a reliable tool for supervised classification of functional data.

APPENDIX A. ADDITIONAL RESULTS

This section summarizes the performance metrics and additional statistical measures for each classifier, evaluated for both experimental and real datasets.

A.1 Statistics for synthetic data

Here, we use 100-fold CV. Tables [A1](#) to [A6](#) report the per-fold accuracy summaries for the six synthetic experiments.

Table A1. Statistics of Experiment 1 (stratified 100-fold outer-CV, with validation features recomputed against training folds only).

Method	Minimum	1st quartile	Median	Mean	3rd quartile	Maximum	NAs
eeh_lda	0.00	0.25	0.50	0.50	0.75	1.00	0.00
eeh_qda	0.00	0.25	0.50	0.51	0.75	1.00	0.00
eeh_knn	0.00	0.25	0.50	0.46	0.75	1.00	0.00
eeh_svm	0.00	0.25	0.50	0.50	0.75	1.00	0.00
eeh_rf	0.00	0.25	0.50	0.45	0.75	1.00	0.00
eee_lda	0.00	0.25	0.50	0.47	0.75	1.00	0.00
eee_qda	0.00	0.25	0.50	0.46	0.75	1.00	0.00
eee_knn	0.00	0.25	0.50	0.45	0.56	1.00	0.00
eee_svm	0.00	0.25	0.50	0.51	0.75	1.00	0.00
eee_rf	0.00	0.25	0.50	0.46	0.56	1.00	0.00
dd_lda	0.00	0.25	0.50	0.50	0.75	1.00	0.00
dd_qda	0.00	0.25	0.50	0.47	0.50	1.00	0.00
dd_knn	0.00	0.25	0.50	0.45	0.50	1.00	0.00

Table A2. Statistics of Experiment 2 (stratified 100-fold outer-CV, with validation features recomputed against training folds only).

Method	Minimum	1st quartile	Median	Mean	3rd quartile	Maximum	NAs
eeh_lda	0.00	0.25	0.50	0.52	0.75	1.00	0.00
eeh_qda	0.75	1.00	1.00	0.97	1.00	1.00	0.00
eeh_knn	0.50	1.00	1.00	0.96	1.00	1.00	0.00
eeh_svm	0.50	0.94	1.00	0.93	1.00	1.00	0.00
eeh_rf	0.50	1.00	1.00	0.97	1.00	1.00	0.00
eee_lda	0.00	0.25	0.50	0.48	0.75	1.00	0.00
eee_qda	0.75	1.00	1.00	0.97	1.00	1.00	0.00
eee_knn	0.50	1.00	1.00	0.96	1.00	1.00	0.00
eee_svm	0.50	1.00	1.00	0.94	1.00	1.00	0.00
eee_rf	0.50	1.00	1.00	0.97	1.00	1.00	0.00
dd_lda	0.50	0.75	1.00	0.86	1.00	1.00	0.00
dd_qda	0.50	1.00	1.00	0.95	1.00	1.00	0.00
dd_knn	0.50	1.00	1.00	0.94	1.00	1.00	0.00

Table A3. Statistics of Experiment 3 (stratified 100-fold outer-CV, with validation features recomputed against training folds only).

Method	Minimum	1st quartile	Median	Mean	3rd quartile	Maximum	NAs
eeh_lda	0.75	1.00	1.00	0.99	1.00	1.00	0.00
eeh_qda	0.75	1.00	1.00	0.99	1.00	1.00	0.00
eeh_knn	0.75	1.00	1.00	0.99	1.00	1.00	0.00
eeh_svm	0.75	1.00	1.00	0.99	1.00	1.00	0.00
eeh_rf	0.75	1.00	1.00	0.99	1.00	1.00	0.00
eee_lda	0.75	1.00	1.00	0.99	1.00	1.00	0.00
eee_qda	0.75	1.00	1.00	0.99	1.00	1.00	0.00
eee_knn	0.75	1.00	1.00	0.99	1.00	1.00	0.00
eee_svm	0.75	1.00	1.00	0.99	1.00	1.00	0.00
eee_rf	0.75	1.00	1.00	0.99	1.00	1.00	0.00
dd_lda	0.75	1.00	1.00	0.99	1.00	1.00	0.00
dd_qda	0.75	1.00	1.00	0.99	1.00	1.00	0.00
dd_knn	0.75	1.00	1.00	0.99	1.00	1.00	0.00

Table A4. Statistics of Experiment 4 (stratified 100-fold outer-CV, with validation features recomputed against training folds only).

Method	Minimum	1st quartile	Median	Mean	3rd quartile	Maximum	NAs
eeh_lda	0.00	0.50	0.50	0.57	0.75	1.00	0.00
eeh_qda	0.50	1.00	1.00	0.96	1.00	1.00	0.00
eeh_knn	0.50	1.00	1.00	0.95	1.00	1.00	0.00
eeh_svm	0.50	0.94	1.00	0.93	1.00	1.00	0.00
eeh_rf	0.50	1.00	1.00	0.96	1.00	1.00	0.00
eee_lda	0.00	0.25	0.50	0.51	0.75	1.00	0.00
eee_qda	0.50	1.00	1.00	0.96	1.00	1.00	0.00
eee_knn	0.50	1.00	1.00	0.95	1.00	1.00	0.00
eee_svm	0.50	0.94	1.00	0.94	1.00	1.00	0.00
eee_rf	0.50	1.00	1.00	0.95	1.00	1.00	0.00
dd_lda	0.00	0.50	0.75	0.69	0.75	1.00	0.00
dd_qda	0.50	0.75	1.00	0.91	1.00	1.00	0.00
dd_knn	0.50	0.75	1.00	0.93	1.00	1.00	0.00

Table A5. Statistics of Experiment 5 (stratified 100-fold outer-CV, with validation features recomputed against training folds only).

Method	Minimum	1st quartile	Median	Mean	3rd quartile	Maximum	NAs
eeh_lda	0.00	0.25	0.50	0.53	0.75	1.00	0.00
eeh_qda	0.50	1.00	1.00	0.96	1.00	1.00	0.00
eeh_knn	0.50	1.00	1.00	0.96	1.00	1.00	0.00
eeh_svm	0.25	0.75	1.00	0.89	1.00	1.00	0.00
eeh_rf	0.50	1.00	1.00	0.95	1.00	1.00	0.00
eee_lda	0.00	0.44	0.50	0.54	0.75	1.00	0.00
eee_qda	0.50	1.00	1.00	0.96	1.00	1.00	0.00
eee_knn	0.50	1.00	1.00	0.97	1.00	1.00	0.00
eee_svm	0.25	0.94	1.00	0.93	1.00	1.00	0.00
eee_rf	0.50	1.00	1.00	0.94	1.00	1.00	0.00
dd_lda	0.50	0.75	1.00	0.93	1.00	1.00	0.00
dd_qda	0.75	1.00	1.00	0.97	1.00	1.00	0.00
dd_knn	0.75	1.00	1.00	0.96	1.00	1.00	0.00

Table A6. Statistics of Experiment 6 (stratified 100-fold outer-CV, with validation features recomputed against training folds only).

Method	Minimum	1st quartile	Median	Mean	3rd quartile	Maximum	NAs
eeh_lda	0.50	0.75	1.00	0.91	1.00	1.00	0.00
eeh_qda	0.50	0.75	1.00	0.91	1.00	1.00	0.00
eeh_knn	0.50	0.75	1.00	0.91	1.00	1.00	0.00
eeh_svm	0.50	0.75	1.00	0.92	1.00	1.00	0.00
eeh_rf	0.25	0.75	1.00	0.89	1.00	1.00	0.00
eee_lda	0.50	0.75	1.00	0.91	1.00	1.00	0.00
eee_qda	0.50	0.75	1.00	0.91	1.00	1.00	0.00
eee_knn	0.50	0.75	1.00	0.91	1.00	1.00	0.00
eee_svm	0.50	0.75	1.00	0.92	1.00	1.00	0.00
eee_rf	0.00	0.75	1.00	0.90	1.00	1.00	0.00
dd_lda	0.50	0.75	1.00	0.91	1.00	1.00	0.00
dd_qda	0.50	0.75	1.00	0.91	1.00	1.00	0.00
dd_knn	0.50	0.75	1.00	0.91	1.00	1.00	0.00

A.2 Statistics for real data

Tables A7, A8, and A9 report the per-fold accuracy summaries for the Berkeley growth study, the Tecator dataset, and the MCO dataset.

Table A7. Statistics of Berkeley growth study data (stratified 20-fold outer-CV, with validation features recomputed against training folds only).

Method	Minimum	1st quartile	Median	Mean	3rd quartile	Maximum	NAs
eeh_lda	0.40	0.60	0.68	0.72	0.85	1.00	0.00
eeh_qda	0.50	0.71	0.80	0.81	1.00	1.00	0.00
eeh_knn	0.40	0.50	0.80	0.73	0.85	1.00	0.00
eeh_svm	0.40	0.57	0.68	0.69	0.80	1.00	0.00
eeh_rf	0.40	0.57	0.80	0.72	0.80	1.00	0.00
eee_lda	0.00	0.50	0.60	0.55	0.60	0.80	0.00
eee_qda	0.40	0.60	0.78	0.80	1.00	1.00	0.00
eee_knn	0.60	0.75	0.80	0.79	0.80	1.00	0.00
eee_svm	0.50	0.60	0.60	0.66	0.80	0.80	0.00
eee_rf	0.20	0.47	0.75	0.67	0.80	1.00	0.00
dd_lda	0.50	0.60	0.78	0.74	0.85	1.00	0.00
dd_qda	0.50	0.60	0.78	0.73	0.80	1.00	0.00
dd_knn	0.20	0.50	0.75	0.68	0.80	1.00	0.00

Table A8. Statistics of Tecator data (stratified 50-fold outer-CV, with validation features recomputed against training folds only).

Method	Minimum	1st quartile	Median	Mean	3rd quartile	Maximum	NAs
eeh_lda	0.25	0.50	0.71	0.68	0.80	1.00	0.00
eeh_qda	0.00	0.50	0.63	0.64	0.80	1.00	0.00
eeh_knn	0.25	0.50	0.63	0.67	0.80	1.00	0.00
eeh_svm	0.25	0.50	0.67	0.66	0.80	1.00	0.00
eeh_rf	0.25	0.60	0.75	0.67	0.80	1.00	0.00
eee_lda	0.25	0.50	0.67	0.68	0.80	1.00	0.00
eee_qda	0.00	0.50	0.63	0.63	0.80	1.00	0.00
eee_knn	0.25	0.50	0.67	0.66	0.80	1.00	0.00
eee_svm	0.25	0.50	0.71	0.66	0.80	1.00	0.00
eee_rf	0.00	0.50	0.75	0.65	0.80	1.00	0.00
dd_lda	0.00	0.60	0.67	0.65	0.75	1.00	0.00
dd_qda	0.00	0.60	0.67	0.65	0.75	1.00	0.00
dd_knn	0.25	0.50	0.63	0.66	0.80	1.00	0.00

Table A9. Statistics of MCO data (stratified 30-fold outer-CV, with validation features recomputed against training folds only).

Method	Minimum	1st quartile	Median	Mean	3rd quartile	Maximum	NAs
eeh_lda	0.00	0.38	0.67	0.62	1.00	1.00	0.00
eeh_qda	0.00	0.50	0.67	0.63	1.00	1.00	0.00
eeh_knn	0.00	0.50	0.67	0.62	0.94	1.00	0.00
eeh_svm	0.00	0.33	0.67	0.61	0.94	1.00	0.00
eeh_rf	0.00	0.50	0.67	0.62	0.75	1.00	0.00
eee_lda	0.00	0.50	0.67	0.60	0.67	1.00	0.00
eee_qda	0.00	0.50	0.67	0.65	0.94	1.00	0.00
eee_knn	0.00	0.33	0.50	0.57	0.94	1.00	0.00
eee_svm	0.00	0.33	0.50	0.57	0.67	1.00	0.00
eee_rf	0.00	0.33	0.67	0.58	0.94	1.00	0.00
dd_lda	0.00	0.33	0.58	0.60	1.00	1.00	0.00
dd_qda	0.00	0.33	0.67	0.62	1.00	1.00	0.00
dd_knn	0.25	0.50	0.67	0.70	1.00	1.00	0.00

STATEMENTS

Author contributions

Conceptualization: F. Zuluaga, H. Laniado. Data curation: A. Carvajal, A. Gomez, C. Lesmes. Formal analysis: F. Zuluaga, H. Laniado. Investigation: A. Carvajal, A. Gomez, C. Lesmes, F. Zuluaga, H. Laniado. Methodology: C. Lesmes, F. Zuluaga, H. Laniado. Software: A. Carvajal, A. Gomez, C. Lesmes. Supervision: C. Lesmes, F. Zuluaga, H. Laniado. Validation: A. Carvajal, A. Gomez, C. Lesmes, F. Zuluaga, H. Laniado. Visualization: A. Gomez, C. Lesmes. Writing of the original draft: C. Lesmes. Writing, reviewing, and editing: A. Carvajal, C. Lesmes. All authors have read and agreed to the published version of the article.

Conflicts of interest

The authors declare no conflict of interest.

Data and code availability

The code used to reproduce the analyses and results presented in this article is publicly available in the GitHub repository github.com/clesmesr/EE-classifier. This repository contains all scripts, data processing workflows, and documentation required to ensure full reproducibility of the study's findings.

Declaration on the use of artificial intelligence (AI) technologies

The authors declare that no generative AI was used in the preparation of this article.

Funding

This research was funded by the project “*Design of new statistical techniques for the classification of high dimensional data*” with number 1119-11190072021 from Universidad EAFIT, Medellín, Colombia.

Open access statement

The Chilean Journal of Statistics (ChJS) is an open-access journal. Articles published in the ChJS are distributed under the terms of the Creative Commons Attribution-NonCommercial-ShareAlike 4.0 License, which permits others to remix, adapt, and build upon the material for non-commercial purposes, provided that appropriate credit is given and derivative works are licensed under identical terms.

REFERENCES

- [1] Ramsay, J.O., 1982. When the data are functions. *Psychometrika*, 47, 379–396. DOI: [10.1007/BF02293704](https://doi.org/10.1007/BF02293704).
- [2] Yao, F., Müller, H.G., and Wang, J.L., 2005. Functional data analysis for sparse longitudinal data. *Journal of the American Statistical Association*, 100, 577–590. DOI: [10.1198/016214504000001745](https://doi.org/10.1198/016214504000001745).

- [3] James, G.M. and Hastie, T.J., 2001. Functional linear discriminant analysis for irregularly sampled curves. *Journal of the Royal Statistical Society B*, 63, 533–550. DOI: [10.1111/1467-9868.00297](https://doi.org/10.1111/1467-9868.00297).
- [4] Li, P.L. and Chiou, J.M., 2011. Identifying cluster number for subspace projected functional data clustering. *Computational Statistics and Data Analysis*, 55, 2090–2103. DOI: [10.1016/j.csda.2011.01.001](https://doi.org/10.1016/j.csda.2011.01.001).
- [5] Horváth, L. and Kokoszka, P., 2012. *Inference for Functional Data with Applications*. Springer, New York, NY, US. DOI: [10.1007/978-1-4614-3655-3](https://doi.org/10.1007/978-1-4614-3655-3).
- [6] Calle-Saldarriaga, A., Laniado, H., Zuluaga, F., and Leiva, V., 2021. Homogeneity tests for functional data based on depth-depth plots with chemical applications. *Chemometrics and Intelligent Laboratory Systems*, 219, 104420. DOI: [10.1016/j.chemolab.2021.104420](https://doi.org/10.1016/j.chemolab.2021.104420).
- [7] Duenas, M. and Giraldo, R., 2022. Multivariate spatial prediction based on Andrews curves and functional geostatistics. *Chilean Journal of Statistics*, 13, 3–16. DOI: [10.32372/ChJS.13-01-01](https://doi.org/10.32372/ChJS.13-01-01).
- [8] Almanjahie, I.M., Attouch, M.K., Fetitah, O., and Louhab, H., 2020. Robust kernel regression estimator of the scale parameter for functional ergodic data with applications. *Chilean Journal of Statistics*, 11, 73–93. URL: soche.cl/chjs/volumes/11/ChJS-11-02-01.
- [9] Hmood, M.Y. and Talib, H.R., 2024. A method of multi-dimensional variable selection for additive partial linear models. *Chilean Journal of Statistics*, 15, 126–146. DOI: [10.32372/ChJS.15-02-02](https://doi.org/10.32372/ChJS.15-02-02).
- [10] López-Pintado, S. and Romo, J., 2005. Depth-based classification for functional data. Technical Report 05-16, Universidad Carlos III de Madrid, Departamento de Estadística, Madrid, Spain. URL: hdl.handle.net/10016/231.
- [11] Li, J., Cuesta-Albertos, J.A., and Liu, R.Y., 2012. DD-classifier: Nonparametric classification procedure based on DD-Plot. *Journal of the American Statistical Association*, 107, 737–753. DOI: [10.1080/01621459.2012.688462](https://doi.org/10.1080/01621459.2012.688462).
- [12] Cuesta-Albertos, J.A., Febrero-Bande, M., and Oviedo de la Fuente, M., 2017. The DD^G -classifier in the functional setting. *TEST*, 26, 119–142. DOI: [10.1007/s11749-016-0502-6](https://doi.org/10.1007/s11749-016-0502-6).
- [13] Franco-Pereira, A.M., Lillo, R.E., and Romo, J., 2011. Extremality for functional data. In: Ferraty, F. (Ed.). *Recent Advances in Functional Data Analysis and Related Topics*. Springer, Heidelberg, Germany, pp. 131–134. DOI: [10.1007/978-3-7908-2736-1_20](https://doi.org/10.1007/978-3-7908-2736-1_20).
- [14] Arribas-Gil, A. and Romo, J., 2014. Shape outlier detection and visualization for functional data: the outliergram. *Biostatistics*, 15, 603–619. DOI: [10.1093/biostatistics/kxu006](https://doi.org/10.1093/biostatistics/kxu006).
- [15] Martin-Barragan, B., Lillo, R.E., and Romo, J., 2016. Functional boxplots based on epigraphs and hypographs. *Journal of Applied Statistics*, 43, 1088–1103. DOI: [10.1080/02664763.2015.1092108](https://doi.org/10.1080/02664763.2015.1092108).
- [16] Sun, Y. and Genton, M.G., 2011. Functional boxplots. *Journal of Computational and Graphical Statistics*, 20, 316–334. DOI: [10.1198/jcgs.2011.09224](https://doi.org/10.1198/jcgs.2011.09224).
- [17] Franco-Pereira, A.M. and Lillo, R.E., 2020. Rank tests for functional data based on the epigraph, the hypograph and associated graphical representations. *Advances in Data Analysis and Classification*, 14, 651–676. DOI: [10.1007/s11634-019-00380-9](https://doi.org/10.1007/s11634-019-00380-9).
- [18] Pulido, B., Franco-Pereira, A.M., and Lillo, R.E., 2023. A fast epigraph and hypograph-based approach for clustering functional data. *Statistics and Computing*, 33, 36. DOI: [10.1007/s11222-023-10213-7](https://doi.org/10.1007/s11222-023-10213-7).

- [19] Tuddenham, R.D. and Snyder, M.M., 1954. Physical growth of California boys and girls from birth to eighteen years. *Publications in Child Development*, 1, 183–364, University of California, Berkeley, CA, US. URL: psycnet.apa.org/record/1956-00215-001.
- [20] Ferraty, F. and Vieu, P., 2006. *Nonparametric Functional Data Analysis: Theory and Practice*. Springer, New York, NY, US. DOI: [10.1007/0-387-36620-2](https://doi.org/10.1007/0-387-36620-2).
- [21] Ruiz-Meana, M., Garcia-Dorado, D., Pina, P., Inserte, J., Agullo, L., and Soler-Soler, J., 2003. Cariporide preserves mitochondrial proton gradient and delays ATP depletion in cardiomyocytes during ischemic conditions. *American Journal of Physiology – Heart and Circulatory Physiology*, 285, H999–H1006. DOI: [10.1152/ajpheart.00035.2003](https://doi.org/10.1152/ajpheart.00035.2003).
- [22] Liu, R.Y., Parelius, J.M., and Singh, K., 1999. Multivariate analysis by data depth: descriptive statistics, graphics and inference. *The Annals of Statistics*, 27, 783–858. DOI: [10.1214/aos/1018031260](https://doi.org/10.1214/aos/1018031260).
- [23] Fraiman, R. and Muniz, G., 2001. Trimmed means for functional data. *TEST*, 10, 419–440. DOI: [10.1007/BF02595706](https://doi.org/10.1007/BF02595706).
- [24] Cuevas, A., Febrero, M., and Fraiman, R., 2007. Robust estimation and classification for functional data via projection-based depth notions. *Computational Statistics*, 22, 481–496. DOI: [10.1007/s00180-007-0053-0](https://doi.org/10.1007/s00180-007-0053-0).
- [25] R Core Team, 2022. *R: A Language and Environment for Statistical Computing*. R Foundation for Statistical Computing. Vienna, Austria. URL: www.R-project.org.
- [26] Venables, W.N. and Ripley, B.D., 2002. *Modern Applied Statistics with S*. Springer, New York, NY, US. DOI: [10.1007/978-0-387-21706-2](https://doi.org/10.1007/978-0-387-21706-2).
- [27] López-Pintado, S. and Romo, J., 2011. A half-region depth for functional data. *Computational Statistics and Data Analysis*, 55, 1679–1695. DOI: [10.1016/j.csda.2010.10.024](https://doi.org/10.1016/j.csda.2010.10.024).
- [28] Ieva, F., Paganoni, A.M., Romo, J., and Tarabelloni, N., 2019. roahd Package: Robust analysis of high dimensional data. *R Journal*, 11, 291–307. DOI: [10.32614/RJ-2019-032](https://doi.org/10.32614/RJ-2019-032).
- [29] Kuhn, M., 2008. Building predictive models in R using the caret Package. *Journal of Statistical Software*, 28, 1–26. DOI: [10.18637/jss.v028.i05](https://doi.org/10.18637/jss.v028.i05).
- [30] Febrero-Bande, M. and Oviedo de la Fuente, M., 2012. Statistical computing in functional data analysis: the R package fda.usc. *Journal of Statistical Software*, 51, 1–28. DOI: [10.18637/jss.v051.i04](https://doi.org/10.18637/jss.v051.i04).

Disclaimer/publisher’s note: The views, opinions, data, and information presented in all publications of the ChJS are solely those of the individual authors and contributors, and do not necessarily reflect the views of the journal or its editors. The journal and its editors assume no responsibility or liability for any harm to people or property resulting from the use of ideas, methods, instructions, or products mentioned in the content.