



國立臺灣大學理學院統計與數據科學研究所

碩士論文

Institute of Statistical and Data Science

College of Science

National Taiwan University

Master's Thesis

使用資料深度來識別各種函數型資料的變化

Identification of various functional changes with the aid of
data depth

吳維哲

Wei-Che Wu

指導教授：陳裕庭 博士

Advisor: Yu-Ting Chen, Ph.D.

中華民國 113 年 7 月

July, 2024



Acknowledgements

研究是一條充滿挑戰的道路，經過重重困難，翻越一座座山巒後，迎來的不會是一片舒適的草原，而是下一座巍峨的高山。然而，我能夠在這不長不短的兩年中完成這篇論文，心中充滿了對許多人的感激之情。在此，我要誠心向所有曾經幫助過我的人表達我的謝意。

從最初的文獻閱讀、研讀及探索，到後來提出研究想法、修改思路及撰寫論文，這兩年內的每一步都離不開我的指導教授陳裕庭老師的悉心指導。最開始時，我在閱讀論文和理解作者意圖上感到萬分困難，是陳老師引導著我，甚至陪同我一起學習；在我提出自己的研究想法後，他不斷提供可行的方向和修改建議；在論文撰寫過程中，從初期拙劣的表達和模糊的文意到本篇論文的完成，都離不開他的耐心指導。甚至在口試過程中，當我表達不清或面對口試委員提出的挑戰性問題時，他的即時幫助使我順利通過口試。這些過程中的每一個環節若沒有他的支持，我都難以克服。在此，向陳裕庭老師致以最崇高的感謝。

我也要感謝口試當日撥冗前來的兩位口試委員楊鈞濤老師及李百齡老師，感謝他們提供了改進方法的建議，提出了研究想法中的短處且對論文內容如何修正的寶貴意見。有了你們的指導，這篇論文才能更加完善和嚴謹。

在此，我還要感謝統計所的各位老師，兩年來扎實的學習讓我一個非本科系的學生學到了該領域的專業知識，也使我能夠自信地與他人分享這兩年來的學習



成果。同時，我也學會了如何思考問題並逐步解決面臨的難題。此外，我還要感謝統計所同學們，他們提供了許多解決問題的創新思路，甚至在期限將至前也願意幫我進行更多的模擬測試使這篇論文的内容更加充實和有說服力。以及特別感謝所上的行政人員，他們辛勤地照顧學生，使我們能專心研究而不必擔心各種繁瑣的雜務。由衷的感謝所各位提供的協助。

最後，我要誠摯地感謝我的家人。能進入統計所學習，多虧了家人的支持。在第一年未能如願考上後，他們依然支持我在家全心全意備考和轉考統計所。每當我遇到挫折和挑戰時，父母總是我情緒的最佳紓解管道，並提供實質性的建議使我能夠克服每一個難關。同時，我也感謝在北部學習的妹妹，每次與她的見面都是我這兩年研究生活中短暫的無憂時光。此外，我還要感謝姑姑奶奶及其他親戚們無微不至的照料，讓我能全身心投入研究而無後顧之憂。儘管在研究上，家人能提供的專業幫助有限，但正是他們對我的信任，讓我有完成研究的勇氣。我難以用言語完整表達我的感激，我只想說，謝謝你們的陪伴。



摘要

在函數型資料中的轉折點檢測方法引起了廣泛關注並持續發展，從早期的“CUSUM”衍生到越來越複雜的損失函數。而先前的方法往往需仰賴動差估計量，對於函數型資料來說既耗時且實用性也較低。為了提高研究方法的可行性，本篇提出利用資料深度來建立一個統計量並將其應用在判定轉折點發生位置，最後再結合樣本分割的方法來確認數據中是否確實發生變化。文章後續的模擬也演示了我們方法的可行性，並突顯了不同參數設置的影響。且在文章的最後，我們會針對研究結果做簡單總結並提出一些可改進的方向。

關鍵字：函數型資料、資料深度、轉折點分析、CUSUM、樣本分割





Abstract

Methods for change-point detection in functional data have garnered significant attention and continue to evolve, from the early derivation of “CUSUM” to the emergence of increasingly complex loss functions. However, previous methods often rely on moment estimators which may be time-consuming and impractical for functional data, especially in estimating higher-order moments. To enhance the feasibility, we propose utilizing data depth to establish a statistical measure for identifying change point locations and combine this statistic with the sample splitting method to confirm whether a change has truly occurred in the data. The simulation results demonstrate the feasibility of our method and highlight the impact of different parameter settings. In the conclusion, we provide a summary of our findings and suggest potential directions for future improvements.

Keywords: Functional Data, Data Depth, Change-Point Analysis, CUSUM, Sample Splitting





Contents

	Page
Acknowledgements	i
摘要	iii
Abstract	v
Contents	vii
Chapter 1 Introduction	1
Chapter 2 Data Depth and Methodology	5
2.1 Review of Functional Data Depth	6
2.1.1 FM-depth	6
2.1.2 Random projection depth	7
2.1.3 Band depth and modified band depth	8
2.2 Method Proposed	10
2.3 Remarks on proposed method	15
Chapter 3 Simulation Study	17
3.1 Simulation setting	17
3.2 Comparison	19
3.3 Window Size Decision	25

Chapter 4 Summary and Discussion

References



29

31

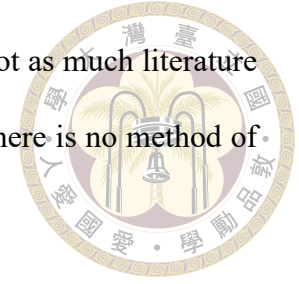


Chapter 1 Introduction

Change point analysis is a subject with a long history and still an active area of research, primarily utilized across a variety of fields such as market stock volatility [Andreou and Ghysels, 2002], engineering quality control and fault detection [Lai, 1995], identification of hazardous factors within genetic sequences [Muggeo and Adelfio, 2011], short-term climate variations [Beaulieu et al., 2012], and even image data [Horváth and Hušková, 2012]. The cases above focus on simple forms of data. With the progress of technology in sampling techniques and data storage methods, the forms of data have become increasingly complex, with high-dimensional and even functional data formats gradually being explored and becoming more common. As a result, datasets in the form of functional data are more common, such as different temperatures in different locations of weather stations, growth curves for boys and girls [Ramsay and Silverman, 2005] and fMRI scan data [Aston et al., 2017] etc. With the increasing number of these datasets, there has been a subsequent emergence and growing importance of analysis methods for functional data. In particular, functional change point analysis is one of the classic and increasingly popular research topics.

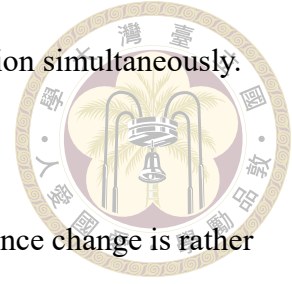
Talking about research on functional change point analysis, it's an issue trying to spot the moments when some significant shifts occurred in the data. In the case of mean

change, numerous methods have been developed, whereas there is not as much literature on covariance change. Additionally, to the best of our knowledge, there is no method of exploring distributional differences currently.



The location shift, or alternatively, change in mean function represents the average behavior of the function over its defined interval. A significant change in the mean results in an obvious deviation in the entire function's location. From best of our knowledge, the earliest study in the case of the mean function has been proposed by [Berkes et al. \[2009\]](#), a statistic generated by integrating the concepts of Functional Principal Component Analysis (FPCA) with CUSUM has also been developed, and they also show the null distribution converges asymptotically to square of a Brownian bridge under the hypothesis H_0 , which indicates a scenario where no change has occurred. Additionally, an asymptotic theory for an estimator constructed by estimated scores of detecting mean change has also been developed, and the finite sample behavior for this estimator has been shown in [Aue et al. \[2009\]](#). However, the assumption of independence for functional data is often too strong. To address the issue of data dependency, [Hörmann and Kokoszka \[2010\]](#) and [Aston and Kirch \[2012\]](#) construct a statistical measure by considering a long-run covariance matrix instead of the ordinary covariance matrix in the proposed statistics under the AMOC (at most one change or no change occurred in our samples) and epidemic change scenarios, respectively. All the methods mentioned above used functional principal component analysis to reduce dimension without considering the information lost by the step. [Aue et al. \[2018\]](#) introduced a method for detecting mean change but without a dimension reduction step. The methods mentioned above focus on a single parameter variable. On the other hand, [Gromenko et al. \[2017\]](#) proposed a test statistics combined with additional weight

to deal with the case of observations consist of time and spatial location simultaneously.



Compared to the literature for mean change, research on covariance change is rather less extensive. To the best of our knowledge, the first skill to address the detection of functional covariance change has been proposed by [Jarušková \[2013\]](#), utilizing FPCA to identify covariance change. Their statistic also follows the square of a Brownian bridge under the condition that no change occurred. [Aue et al. \[2020\]](#) proposed another method for detecting covariance change, which uses the fluctuations of sample eigenvalues or traces of the sample covariance matrix. [Dette and Kokot \[2022\]](#) construct a method for detecting relevant differences in covariance by establishing a null hypothesis in which the form states that the distance between covariances is small. Such a method can also be applied to functional time series cases. Methods mentioned above focus on the AMOC model, but research on multiple changes in covariance has also been proposed. For example, [Harris et al. \[2022\]](#) introduced a Multiple Changepoint Isolation (MCI) method with an augmented fused lasso procedure after data projection. Another case worth discussing is the scenario of high-dimensional functions, [Santo and Zhong \[2020\]](#) proposed homogeneity tests for covariance by computing the trace of the sample covariance matrix and some feasible computation methods. To avoid missing information by using functional principal component analysis, [Jiao et al. \[2023\]](#) introduced a method without dimension reduction for a general situation for weakly dependent observations.

However, all the above ideas must be implemented with moment estimators. Therefore, if we aim to detect covariance changes by using CUSUM, it would necessitate considering four arguments of the operator, which would pose numerous difficulties and render

the verification of hypotheses or practical computations infeasible. To circumvent this dilemma, we attempt to employ another statistical measure to achieve our objective. In doing so, we opt to utilize an alternative index known as “data depth” to circumvent complex computation so that we are able to avoid the aforementioned challenges associated with using moment estimation. Additionally, data depth has the advantage of capturing the relative position of data within the distribution.

The remainder of this article is organized as follows: In Section 2, we will introduce some common functional depths. Following that, we will propose another novel statistic that can be applied to change point detection and explain how we arrived at this idea. In Section 3, a brief simulation study will be conducted to examine the performance of our method under various alternative settings.



Chapter 2 Data Depth and Methodology

In this article, we assume that we obtained independent continuous random function sequence, $x_1(t), x_2(t), x_3(t), \dots, x_{n-1}(t), x_n(t)$ belong to $\mathcal{L}^2$ -space with domain $\mathcal{T} = [0, 1]$.

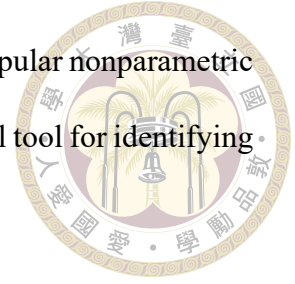
In other words,

$$x_i(t) \in \mathcal{L}^2(\mathcal{T}), \int_{\mathcal{T}} x_i^2(t) dt < \infty, \text{ for } i \in \{1, 2, \dots, n\} \quad (2.1)$$

Furthermore, among these functions, there will be at most one change point, denoted as θ^* such that $x_1(t), \dots, x_{\theta^*}(t)$ remains the same distribution while $x_{\theta^*+1}(t), \dots, x_n(t)$ also share another distribution but distinct from the preceding one. (above well-known for AMOC assumption).

From the basic concept of change-point analysis, it can be elucidated that there will be a certain dissimilarity in the data before and after a certain change occurs. Taking functional data as an example, there may exist differences in the mean function or heterogeneity in scale, or even both simultaneously. Moreover, the difference in shape is also worthy of consideration. From an alternative perspective, this heterogeneity can be conceptualized as the functions preceding a particular change are considered outliers with

respect to the data following that change. Within this framework, a popular nonparametric index, data depth, serves as an appropriate choice as well as a practical tool for identifying outliers.



2.1 Review of Functional Data Depth

To the best of our knowledge, the earliest study in data depth can be traced back to [Liu \[1990\]](#), which starts from the concept of simplices. The core idea of data depth measures is to assess the relative positions of data points within a dataset and identify outliers or provide an overall summary of the dataset.

However, the utilization of data depth for change point detection is not a brand-new idea. According to our comprehension, methods utilizing data depth for change point analysis have been previously introduced by [Chenouri et al. \[2020\]](#), [Ramsay and Chenouri \[2021\]](#) retrospectively, albeit less commonly in the context of functional data [[Ramsay and Chenouri, 2021](#)]. They proposed a nonparametric method that utilizes the CUSUM-based statistic of ranks based on functional depth to detect covariance change. In contrast to [Ramsay and Chenouri \[2021\]](#), this article aims to accurately detect various types of changes. To address the issues we aim to resolve in this article, we would prefer to employ functional data depth intuitively. Hence, a brief review of some commonly used functional data depth measures will be conducted below.

2.1.1 FM-depth

[Fraiman and Muniz \[2001\]](#) introduced a functional data depth, which is called FM-depth. Assume $x_1(t), \dots, x_n(t)$ be some independent stochastic process defined on the in-

terval $[a,b]$, and $[a,b]$ will be set in $[0,1]$, generally. We also define $F_{n,t}(x_i(t))$ as the empirical distribution for curves $x_1(t), \dots, x_n(t)$ ($i = 1, 2, \dots, n$) at a fixed time-point t and $t \in [a,b]$. Then empirical FM-depth is defined as

$$FMD_n(x_i) = \int_a^b D_n(x_i(t)) dt, \quad (2.2)$$

where $D_n(x_i(t)) = 1 - |\frac{1}{2} - F_{n,t}(x_i(t))|$ and $F_{n,t}(x_i(t)) = \frac{1}{n} \sum_{j=1}^n \mathcal{I}(x_j(t) \leq x_i(t))$ and $\mathcal{I}(\cdot)$ is an indicator function.

Due to functional setting, time-points within defined intervals are inherently considered to be of infinite dimension, in (2.2) we consider integral rather than summation consequently.

The underlying intuitive meaning can be explained as first contemplating the centrality of these curves at different time points, followed by extending from a certain time point to encompass the entire interval where the functions are defined.

2.1.2 Random projection depth

Another method for utilizing random projection to construct functional data depth has also been proposed in Cuevas et al. [2007]. Given some continuous functions $x_1, \dots, x_n$ well defined on Hilbert-space $\mathcal{L}^2[0,1]$, and a random direction u (independent from all x_i) will be taken to produce data projection. The result of projecting x_i onto v can be obtained using the inner product in the function space, namely $\langle u, x_i \rangle = \int_0^1 u(t)x_i(t)dt$. Similar rules have also been applied on the projection for the first derivative of observations $\langle u, x_i' \rangle = \int_0^1 u(t)x_i'(t)dt$. Assume we now construct projection on Q direction $\{u_1, \dots, u_Q\}$ (often generated by Gaussian Process in $[0,1]$), then sample version of ran-

dom projection depth obtained by

$$RPD_n(x_i) = \frac{1}{Q} \sum_{q=1}^Q D_n(\langle u_q, x_i \rangle, \langle u_q, x'_i \rangle), \quad (2.3)$$



where $D_n(\cdot)$ refers to a depth function defined on $\mathcal{R}^2$ such as zonoid depth, Mahalanobis depth, etc. Further options and definitions of data depth are introduced in Mosler [2013]. The computation of $D_n(\cdot)$ can utilize previously introduced depth measures or another option called the h-model depth (detail in Cuevas et al. [2006]).

In simple terms, the concept of $RPD_n(x)$ can be explained as that if a curve exhibits lower outlierness among all the original observations, meaning it's closer to the median of all observations, then even after being projected in multiple different directions, although it may have relatively lower depth D_n compared to other observations in some directions, overall it should have a higher depth since we still consider the performance across all projection directions.

2.1.3 Band depth and modified band depth

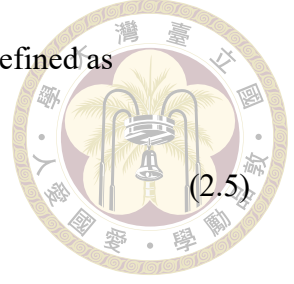
López-Pintado and Romo [2009] established a definition of data depth with the concept of “band” to measure the “centrality” of curves in a functional setting. With a collection of continuous functions $x_1, x_2, \dots, x_n$ defined on the compact interval $\mathcal{T}$, they first constructed the concept of a band which can be denoted as

$$B(x_1, x_2, \dots, x_n) = \{(t, s) : t \in \mathcal{T}, \min_{i=1,2,\dots,n} x_i(t) < s < \max_{i=1,2,\dots,n} x_i(t)\}, \quad (2.4)$$

The above definition represents that at each time point, we identify the range covered by its maximum and minimum values, extending to the entire interval. The union of these

ranges constitutes the “band”. Then, empirical band depth could be defined as

$$BD_{n,J}(x') = \sum_{j=2}^J BD_n^{(j)}(x'), \quad (2.5)$$



where

$$BD_n^{(j)}(x') = \binom{n}{j}^{-1} \sum_{1 \leq a_1 < a_2 < \dots < a_j \leq n} I\{x' \subseteq B(x_{a_1}, x_{a_2}, \dots, x_{a_j})\}. \quad (2.6)$$

An intuitive exposition of band depth of x' can be described as the frequency with a certain curve x' resides in the band formed by other curves within a defined function domain.

However, employing an indicator-based approach to determine the depth of data can be somewhat extreme. For instance, if among n curves $x_1, \dots, x_n$ defined in the interval $[0,1]$, a particular curve x' predominantly aligns with the average of other functions across most intervals but deviating only in the interval $[0,0.01]$ outside the band formed by other curves. A scenario like that may conclude this curve as an outlier when using the previously defined method to calculate its band depth. Hence, an alternative flexible calculation method, termed “modified band depth”, has also been proposed to address this issue.

We first let

$$S_j(x') = S(x'; x_{a_1}, x_{a_2}, \dots, x_{a_j}) = \{t \in \mathcal{T}; \min_{i=1,2,\dots,n} x_i(t) < x' < \max_{i=1,2,\dots,n} x_i(t)\}, \quad (2.7)$$

be the set in the interval $\mathcal{T}$ where x' locate in the band defined with $x_{a_1}, x_{a_2}, \dots, x_{a_j}$, and $1 \leq a_1 < a_2 < \dots < a_j \leq n$. Then “modified band depth” proposed to be constructed with

$$MBD_{n,J}(x') = \sum_{j=2}^J MBD_n^{(j)}(x'), \quad (2.8)$$

where

$$MBD_n^{(j)}x' = \binom{n}{j}^{-1} \sum_{1 \leq a_1 < a_2 < \dots < a_j \leq n} \frac{\Lambda(S_j(x'))}{\Lambda(\mathcal{T})} \quad (2.9)$$

where $\Lambda(\cdot)$ is Lebesgue measure on the interval $\mathcal{T}$.



From the definition above, the disparities in shapes between curves can dramatically influence the magnitude of Band Depth (BD), whereas Modified Band Depth (MBD), in contrast to focusing on the shapes of curves, places greater emphasis on the magnitudes between curves.

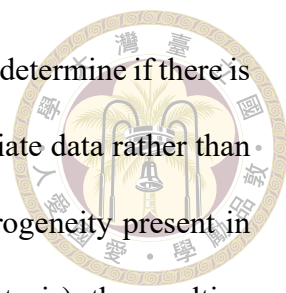
2.2 Method Proposed

With the foundational understanding of functional data depths (FMD, RPD, BD, MBD), we now assume a change occurred at $\theta^* \in \{1, 2, \dots, n\}$. Our original idea stems from a tool for assessing outliers called DD-plot proposed by [Liu et al. \[1999\]](#). Assume $\mathbf{Y} = \{y_1, \dots, y_{n_1}\}$ and $\mathbf{Z} = \{z_1, \dots, z_{n_2}\}$ be random samples from $\mathcal{F}_1$ and $\mathcal{F}_2$ respectively, and these samples can include univariate data, multivariate or even functional. Then a two-dimensional DD-plot can be built with $d(\mathcal{F}_{1n_1}, \mathcal{F}_{2n_2})$ where

$$d(\mathcal{F}_{1n_1}, \mathcal{F}_{2n_2}) = \{(D_{\mathcal{F}_{1n_1}}(x), D_{\mathcal{F}_{2n_2}}(x)), x \in \{\mathbf{Y} \cup \mathbf{Z}\}\} \quad (2.10)$$

$\{\mathcal{F}_{1n_1}, \mathcal{F}_{2n_2}\}$ is the empirical distribution of $\{\mathcal{F}_1, \mathcal{F}_2\}$ and various types of data depth can be used for computing empirical depth $D_{\mathcal{F}_{1n_1}}(\cdot)$. For instance, simplicial depth [\[Liu, 1990\]](#) can be applied to the multivariate case, while modified band depth [\[López-Pintado and Romo, 2009\]](#) is suitable for functional samples.

Upon completion of the DD-plot construction described above, in [Li and Liu \[2004\]](#),



a test statistic constructed from the concept of DD-plot is employed to determine if there is heterogeneity in the distributions of two datasets. Although multivariate data rather than functional data are used, it can be observed that when there is heterogeneity present in the data (including centrality, dispersion, and even skewness and kurtosis), the resulting DD-plot exhibits various specific patterns but share a common characteristic of slightly deviating from the 45-degree line. From the perspective of change point analysis, the existence of heterogeneity can be interpreted as a shift in the data before and after θ^* . Specifically, the data $\{x_1, \dots, x_{\theta^*}\}$ will serve as extreme values relative to $\{x_{\theta^*+1}, \dots, x_n\}$, and vice versa. Then, a significant deviation from the 45-degree line must be expected if a DD-plot is constructed with these two groups segmented by θ^* .

For this reason, the total distances of all points on the DD-plot to the 45-degree line will increase significantly as the heterogeneity in the data becomes stronger. We have developed a statistic to detect whether a change has occurred in our datasets. The statistic can be formulated as:

$$T_0(\theta) = \sum_{i=1}^n |D_{\theta}^B(x_i) - D_{\theta}^A(x_i)| \quad (2.11)$$

where $D_{\theta}^B(x_i)$ represents the data depth of x_i with respect to dataset combined with $\{x_1, x_2, \dots, x_{\theta-1}, x_{\theta}\}$, $D_{\theta}^A(x_i)$ represents the data depth of x_i with respect to the remaining samples $\{x_{\theta+1}, \dots, x_{n-1}, x_n\}$ and θ is our change point predictor.

After fully defining the aforementioned statistic, as previously discussed, we expect that the statistic will be significantly higher at the actual change point than at a location without a change. Consequently, we define the location where this statistic reaches its maximum value as our change point estimator intuitively. Finally, due to the lack of clear understanding regarding the theoretical properties and true distribution of this statistic, we

then employ a permutation test to confirm whether the estimated change point represents an actual change.



Nevertheless, the proposed statistic is not flawless. A noticeable issue arises from the potentially unequal sample sizes of the two segments divided by θ . In that case, an insufficient representative may occur such that the dataset distribution is difficult to capture when the reference group is unequal. This scenario increases the possibility that the estimated change point will fall on a boundary rather than at the actual change when using (2.11). Our simulation tests have also yielded erroneous results in such cases.

To address the aforementioned drawbacks, we note that the change point reflects local behavior in the data [Niu and Zhang, 2012] and incorporates the idea of a sliding window. Continuing from the previous assumption that θ is our change point predictor and dividing the data into two parts $\{x_1, \dots, x_\theta\}$ and $\{x_{\theta+1}, \dots, x_n\}$, we now turn to consider a window segment which used for identifying the change point involves $\lfloor nw \rfloor$ samples before and after θ , noted as $\{x_{\theta-\lfloor nw \rfloor+1}(t), \dots, x_\theta(t)\}$ and $\{x_{\theta+1}(t), \dots, x_{\theta+\lfloor nw \rfloor}(t)\}$, where $\{x_{\theta-\lfloor nw \rfloor+1}(t), \dots, x_\theta(t)\}$ means the set consists of the functions before a change occur, which is often described as the left window, and $\{x_{\theta+1}(t), \dots, x_{\theta+\lfloor nw \rfloor}(t)\}$ represent the remaining functions (right window), where $w \in (0, 0.5)$ is a proper rate. The influence of w will be discussed in the next section.

If $\theta = \theta^*$, the actual change θ^* indicates that data from the left window will be outliers relative to the data in the right window, or data from the right window will be outliers relative to the data in the left window. In the case of a mean change, the data from the two groups act as extreme values for each other. In contrast, only one-sided conformity will occur under a scale change. By integrating the concept of data depth,

we can infer if a sample $x_i(t)$ originates from the left window, the depth of $x_i(t)$ with respect to the right window should be much lower than the depth of $x_i(t)$ with respect to the left window intuitively due to the dissimilarity between two window segments. The calculated significantly lower data depths can subsequently be employed to pinpoint the exact locations of changes within the data.

Based on the above, to accomplish change point detection, we now propose a statistical measure combined with sliding windows which can be defined as

$$T(\theta) = \min\left\{ \sum_{i=\theta-\lfloor nw \rfloor+1}^{\theta} D_L^{\theta}(x_i), \sum_{i=\theta+1}^{\theta+\lfloor nw \rfloor} D_R^{\theta}(x_i) \right\} \quad (2.12)$$

where $D_L^{\theta}(x_i)$ represents the data depth of x_i in the left window with respect to the right window segment and $D_R^{\theta}(x_i)$ is built with same procedure but take reference in the left window segment. Among all candidate $\theta \in \Theta_w = \{x_{\lfloor nw \rfloor+1}, \dots, x_{n-\lfloor nw \rfloor}\}$. We then define

$$T = \min_{\theta \in \Theta_w} T(\theta) \quad (2.13)$$

When the statistic (2.13) reaches a sufficiently low value, we consider that a change has occurred in the data. Similar to the previous method, a permutation test can be used for validation in testing whether an actual change exists. And the most probable location of the change, denoted as $\hat{\theta}^*$, is then defined as

$$\hat{\theta}^* = \underset{\theta \in \Theta_w}{\operatorname{argmin}} T(\theta) = \underset{\theta \in \Theta_w}{\operatorname{argmin}} \min\left\{ \sum_{i=\theta-\lfloor nw \rfloor+1}^{\theta} D_L^{\theta}(x_i), \sum_{i=\theta+1}^{\theta+\lfloor nw \rfloor} D_R^{\theta}(x_i) \right\} \quad (2.14)$$

represents our final change point estimator.

Nevertheless, we opt for sample splitting as an alternative approach due to considerations of computational feasibility. The sample splitting technique is effective for detecting

the number of change points and also estimating their locations precisely [Chong, 2004], [Zou et al., 2020]. In this study, we try to utilize an alternative concept for sample splitting. Initially, we segment the sample functions into odd and even sub-samples noted as $\mathcal{S}_O, \mathcal{S}_E$ where

$$\mathcal{S}_O = \{x_1(t), x_3(t), \dots, x_{2j+1}(t), \dots\} \& \mathcal{S}_E = \{x_2(t), x_4(t), \dots, x_{2j}(t), \dots\}, \forall j = 1, 2, \dots, \lfloor n/2 \rfloor \quad (2.15)$$

Under the assumption of AMOC condition and the existence of a real change, the location identified in the odd and even sub-samples should be approximately the same. Conversely, if the locations of the change points identified in the odd and even sub-samples differ obviously, the guess that a change has occurred in the data should be denied intuitively. A detailed execution summary for the above procedure can be presented as follows.

- (a) Divide n functions into two subset $\mathcal{S}_O, \mathcal{S}_E$ as defined above.
- (b) In each segment group $\mathcal{S}_O, \mathcal{S}_E$, use equation (2.14) to determine the location of change points $\hat{\theta}_O^*$ and $\hat{\theta}_E^*$ in alternative candidate subset Θ_O, Θ_E where $\Theta_O = \{2\lfloor \frac{nw}{2} \rfloor + 1, 2\lfloor \frac{nw}{2} \rfloor + 3, \dots, n - 2\lfloor \frac{nw}{2} \rfloor - 1\}$, and $\Theta_E = \{2\lfloor \frac{nw}{2} \rfloor + 2, \dots, n - 2\lfloor \frac{nw}{2} \rfloor\}$.
- (c) With a proper C , we consider that a change has occurred if $|\hat{\theta}_O^* - \hat{\theta}_E^*| \leq C$. Otherwise, concluding that no change has occurred in the functions should be a more persuasive argument.

In the aforementioned procedure, we first divide the sample into odd and even segments. And then within these two segments, we use (2.14) to identify our change point estimator $\hat{\theta}_O^*$ and $\hat{\theta}_E^*$ in each segment. At last, based on the intuitive idea presented earlier, we determine the authenticity of the change θ^* by assessing whether the difference in

position between the two estimators is sufficiently small by using an appropriate constant C .



As for the criterion C , it could be determined based on our specific requirements. If a change in the data is considered critical and requires prevention, we may increase the value of C to ensure that the change can be detected more easily. Conversely, if we wish to avoid detecting less significant changes, we can reduce the value of C to minimize the possibility of identifying such changes. We will adjust C based on the sample size primarily because the proposed statistic uses a sliding window concept, where each change point candidate has an equal chance of being selected as the final estimator. So that adjust the value of C based on the sample size would be a more reasonable approach.

2.3 Remarks on proposed method

At last, we provide a brief analysis and discussion of the strengths and weaknesses of the aforementioned methods and how various parameters influence the proposed methodology. The first discussion focuses on the advantages and disadvantages of using local samples versus the entire sample. The former approach not only mitigates the issue of the estimator frequently falling on the boundary but also avoids the unstable estimated change point experienced by the CUSUM-based methods due to shifting in the change location. However, compared to the latter approach, parameters such as window size significantly impact the estimation results of the proposed method. Consequently, determining the appropriate window size presents a complex and challenging problem.

Another noteworthy issue is the replacement of the permutation test with the sample splitting method. Although this alternative method significantly enhances computational

efficiency thereby increasing the feasibility of the proposed statistics, it still presents a clear drawback which is a partial loss of test power. This power reduction is primarily due to the division of the sample, as opposed to utilizing the entire dataset for testing, which predictably results in a lower testing power outcome. Besides, sample splitting in dependent cases faces challenges due to the correlation among the data. On the other hand, CUSUM-based methods encounter the issue of selecting an appropriate bandwidth when estimating the long-run covariance.

The remaining issue is on the selection of the rate w . For the rate w , in practical applications, the number of candidates defined by Θ_w in this article increases as the sample size becomes larger, thereby reducing the possibility of each candidate being selected as the final estimator compared to smaller sample sizes. Therefore, we adopt a proportion for setting w both before and after sample splitting, typically employing $w = 0.15$ in this study. The simulation and discussion sections will present the results obtained with different w along with more detailed discussions. Regarding the cutoff C , we also consider that the cutoff should be appropriately adjusted its scale according to the sample size.



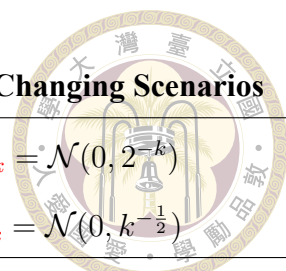
Chapter 3 Simulation Study

The aim of this chapter is to evaluate the performance of the proposed statistic in the AMOC scenario. Comprehensive changing scenarios such as changes in the decay rate, changes in the mean function, changes in the distribution of the random term, and simultaneous fluctuations in both mean function and decay will be complemented to evaluate the behavior of the proposed statistic computed with MBD. We then compare our method with [Ramsay and Chenouri \[2021\]](#) and [Harris et al. \[2022\]](#) to demonstrate the advantage of our method. We will also demonstrate how the window size w mentioned earlier affects the estimation results of the proposed statistic.

3.1 Simulation setting

We set the model without changes as $M_0 : x_i(t) = \mu_0(t) + \sum_{k=1}^{21} \xi_{0ik} \phi_k(t)$, $\forall i \in \{1, 2, \dots, n\}$, ξ_{0ik} generated by $\mathcal{N}(0, k^{-1})$ with gradual decay, $\mu_0(t) = 0.9 + 3t^3(1 - t)$ and $\phi_k(t)$ is the common Fourier basis.

We emulate the simulation setups of [Aue et al. \[2018\]](#) and [Chiou et al. \[2019\]](#), while introducing some minor modifications. We summarize our settings in Table 3.1 and a visualization of a few curves from our settings is presented in Fig 3.1. In models 1 and 2, we set up scenarios where the data undergo fast decay and slow decay respectively. We



Model	Equation	Varying Changing Scenarios
Model 1 (SC1)	$x_i(t) = \mu_0(t) + \sum_{k=1}^{21} \xi_{1ik} \phi_k(t)$	$\xi_{1ik} = \mathcal{N}(0, 2^{-k})$
Model 2 (SC2)	$x_i(t) = \mu_0(t) + \sum_{k=1}^{21} \xi_{2ik} \phi_k(t)$	$\xi_{2ik} = \mathcal{N}(0, k^{-\frac{1}{2}})$
Model 3 (MC1)	$x_i(t) = \mu_1(t) + \sum_{k=1}^{21} \xi_{0ik} \phi_k(t)$	$\mu_1(t) = \mu_0(t) + 0.4 \sin(1 + 10\pi t)$
Model 4 (MC2)	$x_i(t) = \mu_2(t) + \sum_{k=1}^{21} \xi_{0ik} \phi_k(t)$	$\mu_2(t) = 0.8 + 3t^2 - 5t^3$
Model 5 (DC1)	$x_i(t) = \mu_0(t) + \sum_{k=1}^{21} \xi_{ik}^{\mathcal{U}} \phi_k(t)$	$\xi_{ik}^{\mathcal{U}} \in \mathcal{U}(-\frac{6}{k^2}, \frac{6}{k^2})$
Model 6 (DC2)	$x_i(t) = \mu_0(t) + \sum_{k=1}^{21} \xi_{ik}^{\mathcal{L}} \phi_k(t)$	$\xi_{ik}^{\mathcal{L}} \in \mathcal{L}(0, \frac{1}{\sqrt{2k}})$
Model 7 (MSC1)	$x_i(t) = \mu_1(t) + \sum_{k=1}^{21} \xi_{2ik} \phi_k(t)$	contamination in both $\mu_1(t)$ and ξ_{2ik}
Model 8 (MSC2)	$x_i(t) = \mu_2(t) + \sum_{k=1}^{21} \xi_{2ik} \phi_k(t)$	contamination in both $\mu_2(t)$ and ξ_{2ik}

Table 3.1: Model setting

establish a mean shift in shape in model 3 and apply a magnitude fluctuation within model 4. For models 5 and 6, we use the uniform distribution and the Laplace distribution to mimic the scenarios that FPC scores are from different distributions. These distributions are set to have the same mean and variance as M_0 . Finally, in models 7 and 8, we aim to depict more complex forms of change combined slow decay with shape and magnitude respectively. The following discussion will be performed with the data before the change are modeled as M_0 , and the data after the change are modeled according to the aforementioned 8 models. All these scenarios will be set up in 200 simulation runs with proper $w = 0.15$. Simulated datasets combine with independent functions in two sample sizes $n = 200$ and 500 , and the true change points θ^* will be set at the same position $\lfloor 0.3n \rfloor$. Furthermore, we define accuracy as $\#\{|\hat{\theta}^* - \theta^*| \leq C/2\} / \#\{\text{run times}\}$ where $C = \lfloor 0.1n \rfloor$ to facilitate the explanation.

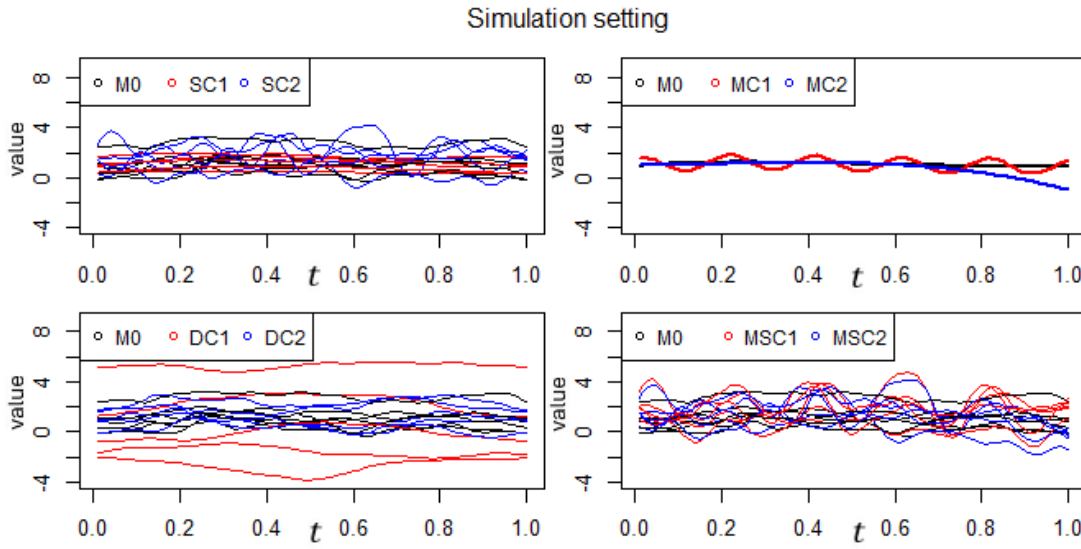


Figure 3.1: Visualization of our setting. For a concise expression, we draw 5 curves in each scenario and only plot the mean function in the mean change case (top right).

3.2 Comparison

In the presentation of simulation results, we begin with a validation of the performance of the statistical measures proposed in this study with $n = 500$. To present the method proposed in this study more clearly, we have selected model 1 and model 4 which exhibit stronger changes, for preliminary demonstration. The left side of Figure 3.2 illustrates that our statistical measures perform well in model 1 and model 4. A conspicuous downward trend at the correct change location ($\theta^* = 150$) can be discovered. The histogram on the right side also demonstrates that our estimator $\hat{\theta}^*$ lands around the precise location where an actual change occurs in the data. Above results validated that our idea is successful.

We then compare the change point locations identified by our method against those captured by [Ramsay and Chenouri \[2021\]](#) (noted as RC) and [Harris et al. \[2022\]](#) (noted as MCI) in sample size $n = 200, 500$. The reason for choosing to compare with [Ramsay and](#)

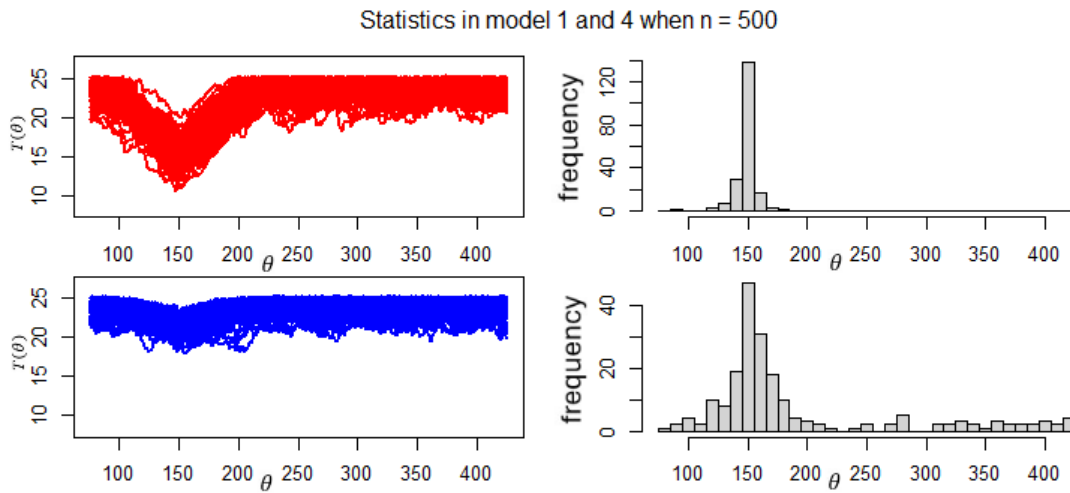


Figure 3.2: The results of the statistics and location under two varying conditions are illustrated. The top of the figure displays the trends in our statistics and our change point estimator in model 1, and the bottom of the figure represents the performance in model 4, and θ^* is set on 150.

Chenouri [2021] is that it utilizes data depth as well but applies it to transform obtained sequences into ranks and execute CUSUM with these ranks to determine change points. This method is primarily used to detect changes in covariance, which is also included in our simulation design. Harris et al. [2022] proposed a method that projects variations both between and within samples and integrates the fused lasso procedure to decide change point candidates, and ultimately applies FDR correction to identify changes in the data. We chose this method primarily because it employs CUSUM within a specific interval after segmentation, indicating that it also relies on local data. Additionally, it can simultaneously detect changes in both the mean and covariance of functional data, which is included in our setting scenarios as well.

We then compare the change point estimator in the above three methods respectively. Since MCI tends to identify multiple change estimators, we select the most significant one—the point with the smallest p-value after FDR correction as the change point identified by MCI. As illustrated in Figure 3.3 and Figure 3.4. Although the most significant estimator of MCI exhibits an overall smaller bias than the other two methods, it tends to overestimate

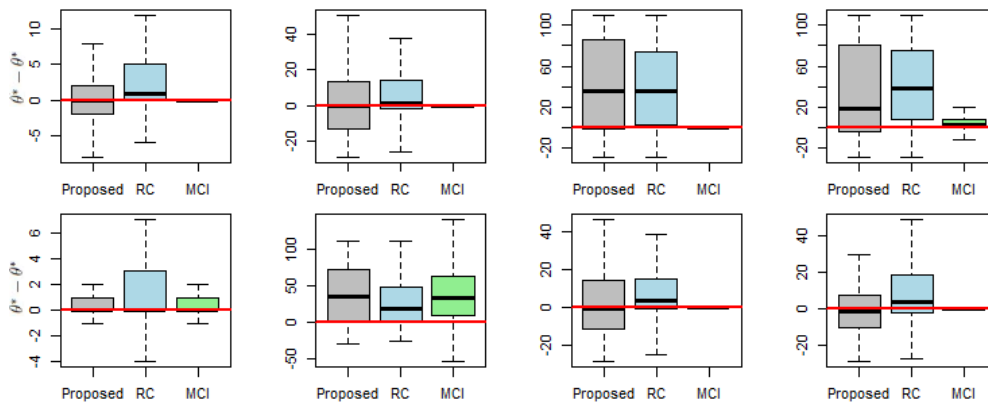


Figure 3.3: The figure above compares estimated bias in 8 models with sample size $n = 200$. The y-axis represents $\hat{\theta}^* - \theta^*$. The first from the top left to the second from the top left toward the bottom right corresponds to models 1,2,..., and 8, respectively.

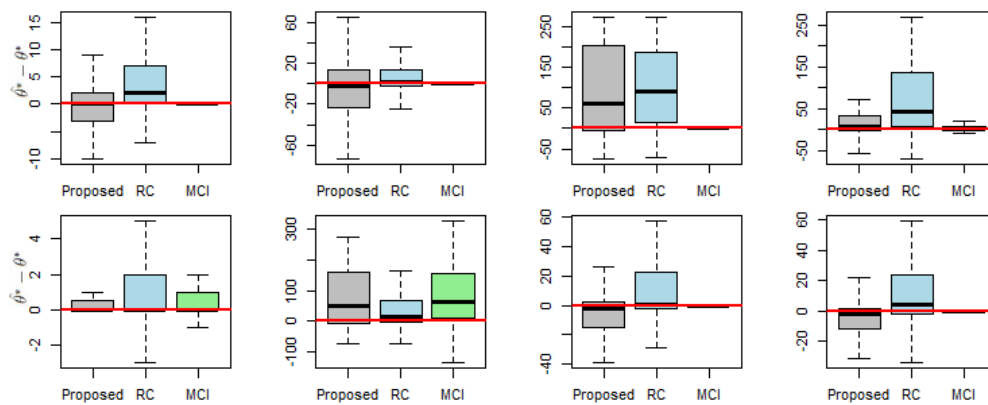
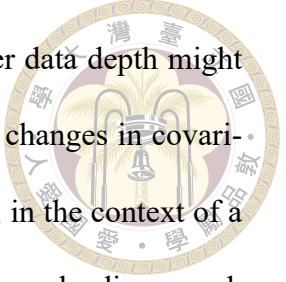


Figure 3.4: The figure above compares estimated bias in 8 models with sample size $n = 500$. The y-axis represents $\hat{\theta}^* - \theta^*$. The first from the top left to the second from the top left toward the bottom right corresponds to models 1,2,..., and 8, respectively.

the number of change points. Therefore, it is challenging to evaluate the performance of this method under the AMOC framework. Consequently, we will primarily focus on discussing the proposed method and the RC method.

In models 1 and 2, our method has better performance than the method in Ramsay and Chenouri [2021] because of CUSUM-based method tends to capture a change in the middle of samples; under the scenario of mean change, the performance of both methods fell short of expectations especially in model 3 which also represents shape change scenario. The primary reason for the weak performance of our method may be attributed

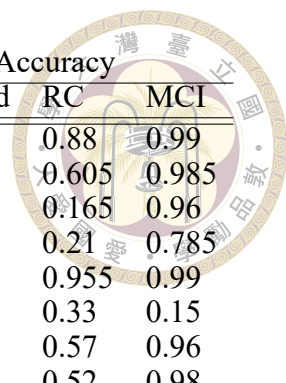
to the instability of MBD to detect shape outliers. Employing another data depth might yield improvements. The RC method is primarily designed to detect changes in covariance, and its insensitivity to mean change is anticipated; nevertheless, in the context of a magnitude shift, i.e. model 4, a significant improvement in our method can be discovered. The primary reason is that our method utilizes more samples as the sample size increases, thereby enhancing the representative of samples within the window. This is an advantage of adjusting the window size proportionally rather than using a fixed number of samples (further discussed in next section). In model 5, the significant disparity between normal and uniform distributions enabled both methods to perform quite well, but in model 6, the resemblance between normal and Laplace distributions, both of which approximate a bell shape, hindered both methods from detecting such changes. Lastly, in models 7 and 8, we obtain similar conclusions to those of models 1 and 2 due to the existence of covariance change.



Index	Power	Size	Accuracy
Proposed	$\frac{\#\{ \theta_0^* - \hat{\theta}_E^* \leq C/2\}}{\#\{run\ times\}}$	$\frac{\#\{ \theta_0^* - \hat{\theta}_E^* \leq C/2\}}{\#\{run\ times\}}$	$\frac{\#\{ \hat{\theta}^* - \theta^* \leq C/2\}}{\#\{run\ times\}}$

Table 3.2: Empirical power, size and accuracy for our method

We now present the performance in change detection for the three methods. The compared index for our method is defined in Table 3.2, and the calculation of empirical power will be discussed in the 8 types of changes mentioned earlier, while size will be calculated under the scenario where no changes have occurred in the data. In RC the empirical power and empirical size are calculated using the 95% quantile of the Brownian bridge derived in [Ramsay and Chenouri \[2021\]](#) and then divided by 200. But for MCI, the empirical power is computed as the number of times one or more changes are detected divided by 200, while the empirical size is computed as the number of times one or more



Sample	Model	Power(Size)			Accuracy		
		Proposed	RC	MCI	Proposed	RC	MCI
$n = 200$	SC1	0.85(0.23)	1(0.12)	0.31(0.97)	0.885	0.88	0.99
	SC2	0.445	0.885	0.14	0.72	0.605	0.985
	MC1	0.26	0.17	0.225	0.175	0.165	0.96
	MC2	0.285	0.26	0.49	0.445	0.21	0.785
	DC1	0.985	1	0.045	0.99	0.955	0.99
	DC2	0.39	0.465	0.245	0.335	0.33	0.15
	MSC1	0.485	0.85	0.175	0.705	0.57	0.96
	MSC2	0.5	0.725	0.39	0.755	0.52	0.98
$n = 500$	SC1	0.935(0.24)	1(0.145)	0.28(0.96)	0.97	0.96	1
	SC2	0.495	1	0.12	0.61	0.85	1
	MC1	0.23	0.185	0.125	0.255	0.205	0.99
	MC2	0.38	0.415	0.59	0.62	0.32	0.91
	DC1	1	1	0.025	1	0.99	0.985
	DC2	0.31	0.66	0.125	0.29	0.495	0.235
	MSC1	0.575	0.985	0.075	0.765	0.78	0.995
	MSC2	0.32	0.98	0.54	0.775	0.75	0.975

Table 3.3: Empirical result comparison in 8 models

changes are detected when there is no change occurred, also divided by 200. The term “accuracy” has been previously defined in Section 3.1. But for MCI, we selected the most significant estimator (the estimator corresponds to the lowest p-value) identified by that for comparison with others, primarily due to MCI tends to identify more change points.

From Table 3.3, we can observe that the method proposed in this article exhibits a slightly inferior performance in empirical power compared to RC overall. Even though an increase in sample size can lead to better results in terms of both accuracy and empirical power for our method, with a more noticeable improvement in accuracy, it still falls short in terms of power when compared to RC. The primary reason for this discrepancy lies in our choice of using the less stable sample splitting approach instead of testing, which is an area where our method could be improved. Despite our method obtaining a slightly inferior performance in terms of empirical power, it generally outperforms RC in accuracy, especially in capturing the precise location of the change. This advantage is even more pronounced in cases involving mean change. Another noteworthy aspect is the performance

of the MCI method, we can observe that MCI has a very low probability of correctly detecting a single change point across the eight scenarios we designed. Even in scenarios where no changes are present, this method frequently detects changes, which can be attributed to its tendency to identify more change points. Nevertheless, when examining accuracy, we find that “the most significant estimator” identified by MCI is generally more precise than the other two methods and falls within our defined standards ($|\hat{\theta}^* - \theta^*| \leq C/2$). Nonetheless, in Model 6 (DC2), all these three methods exhibit poor accuracy due to the weak signal of the change, including MCI which typically performs better overall.

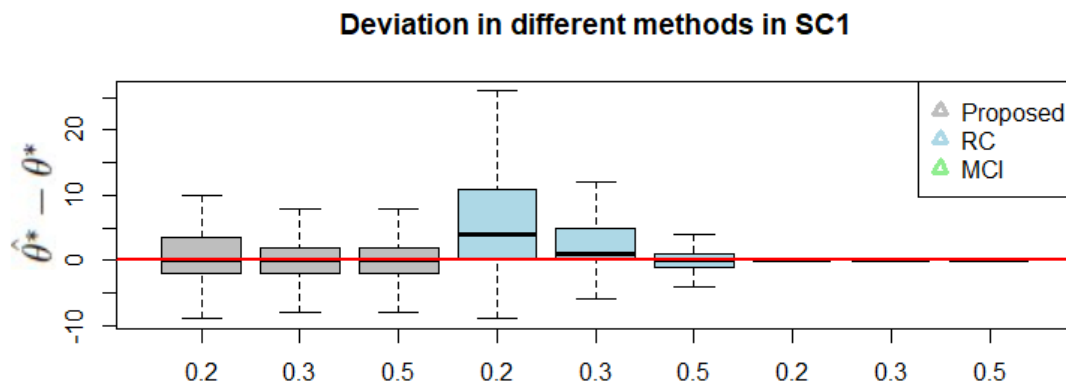
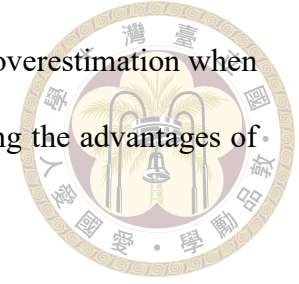


Figure 3.5: The figure illustrates the deviation from the estimated change point to θ^* in three mentioned methods when moving θ^* to three different locations, the result obtained by SC1 with $n = 200$

Another noticeable issue is the performance in different change point locations θ^* . We now set θ^* at $[0.2n], [0.3n], [0.5n]$ in SC1 to compare the accuracy of the change point estimators obtained by three methods. As shown in Figure 3.5, our method are less affected by the variation in the change point location which supports the advantages of using local sample rather than using the entire sample. Same conclusion can be obtained for MCI also identifies change points by examining multiple intervals which represents it also using local data. In contrast, RC uses the entire dataset and relies on CUSUM, tends to

detect changes around the center positions. This leads to noticeable overestimation when design change is located on $\lfloor 0.2n \rfloor$ and $\lfloor 0.3n \rfloor$, further demonstrating the advantages of using local data.



3.3 Window Size Decision

We subsequently engage in a discussion regarding the window size at the conclusion of the simulation. Primarily, we identify two reasonable approaches for selection. The first approach involves utilizing a fixed proportion w , while the second approach entails directly assigning a fixed numerical value w' . Both are reasonable choices. But when the sample size increases or decreases, the number of candidate points also varies with our choice of window size. Nevertheless, with different sample size, we tend to prefer a standard with the same scale. Using w' to implement our method will increase our candidate set such that making us more susceptible to false detection (detailed discussion in Figure 3.6). Therefore, we choose to implement our method by using a fixed proportion w rather than w' due to the insensitivity of w' to sample size and that will allows us to better capture the benefits of an increased sample size.

In Figure 3.6, we can observe that empirical power, size, and accuracy all significantly increase with the rise of w . The empirical power and size even accuracy do not show a significant ascension with further increases in w beyond $w = 0.15$. While in terms of accuracy, it continues to improve as the window size increases as expected because a larger window size better captures the distribution of the samples, meaning the samples within the window are more representative.

Nevertheless, a counterintuitive result is that our empirical size increases with w .

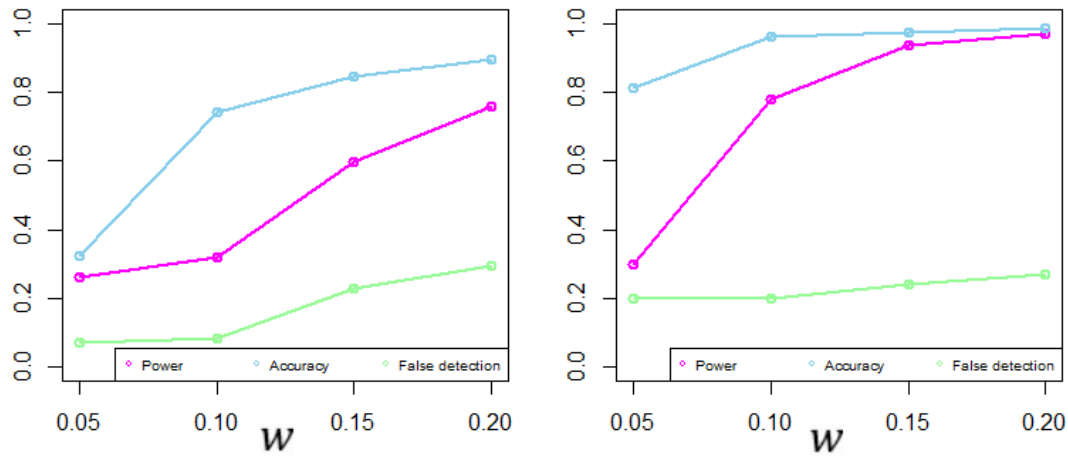


Figure 3.6: Comparison in empirical power, size and accuracy with the same definition in aforementioned section. The left figure illustrates an empirical result of our method in $n = 200$ in SC1 while the right one presents the same result but adjusts sample size to $n = 500$.

The primary reason is that when no change occurs in the data, the minimum value of our test statistic will be randomly distributed within our candidate set Θ_w . As w increases, a smaller candidate set will be obtained, and the possibility of encountering overlapping $\hat{\theta}_O^*$ and $\hat{\theta}_E^*$ increases significantly with w . Namely, the defined criterion $|\hat{\theta}_O^* - \hat{\theta}_E^*|$ in section 2.2 becomes easier to meet regardless of whether a change has actually occurred in the data or not. This indicates that we cannot blindly increase w .

In Figure 3.7, we demonstrate the proposed statistic in some simulation rounds with $n = 200$ in SC1, and a downward trend near θ^* can be observed frequently but a more unstable oscillation occurs when $w \leq 0.1$ than those in $w \geq 0.15$.

In summary, we aim to increase w to enhance the representativeness of our samples for more accurate estimation, but we must also avoid indiscriminately increasing w to prevent false detections. Additionally, increasing w leads to discarding too many samples, or we say, change point candidates. For practical applications, we end up selecting $w = 0.15$ as a reasonable and balanced choice.

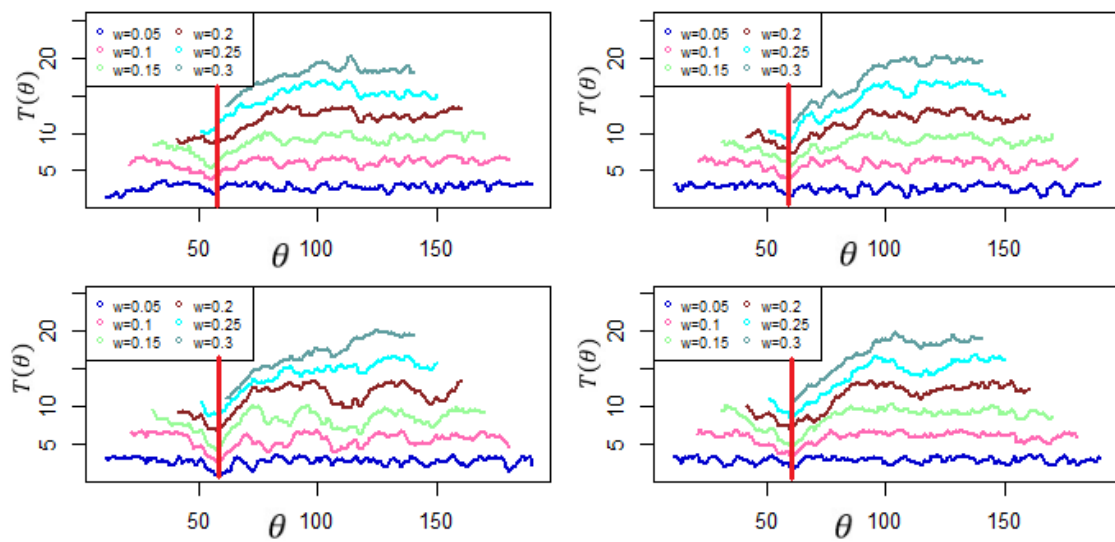


Figure 3.7: The figure illustrates the performance of our statistical measure in different w in SC1 and $n = 200$. The red line represents $\theta^* = 60$. The x-axis represents candidate in Θ_w .





Chapter 4 Summary and Discussion

This article adopts an approach based on data depth and DD-plot, further extending the concept of outlier and linking it to change point detection. The statistical measure is derived using data depth which is commonly used to identify whether the data are outliers. When determining whether a change is actual or not, we apply sample splitting to divide the samples but we assess whether the detected change points in these two samples are approximately the same to validate the detected signal's authenticity. The simulations show the results of our approach and further present the impact of using different window sizes. A brief comparison with the other method is also present our method demonstrates a competitive accuracy in determining the location of changes.

Nevertheless, the statistical measure remains significant room for development. For instance, we have not yet determined the optimal method for selecting w in this article, but have only discussed the pros and cons of different w values. Combining w with our proposed statistical measure using a penalty approach may be a feasible idea. we have not learned detailed information on the distribution of the starting point $T_0(\theta)$ and the final statistical measure $T(\theta)$ under large samples. Understanding the theoretical properties of these statistical measures should enable significant advancements in our method. Additionally, when faced with scenarios such as Model 3 or Model 7, where functional data exhibits behavior often characterized as shape outliers, using Modified Band Depth

(MBD) for computation shows limitations in capturing changes effectively. Exploring alternative depth measures such as Total Variation Depth (TVD)[[Huang and Sun, 2019](#)], which are less influenced by shape outliers, could potentially enhance the performance of the proposed statistical measure. Finally, this study employs the sliding window as the core of the method's construction. Since the sliding window is also a common and feasible approach in online detection, the ideas proposed here may not be limited to offline cases and could potentially perform well in online detection scenarios as well.

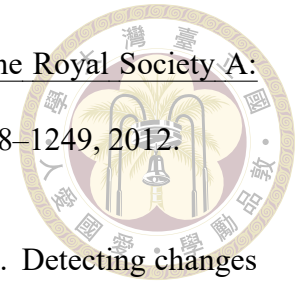




References

- Elena Andreou and Eric Ghysels. Detecting multiple breaks in financial market volatility dynamics. Journal of applied Econometrics, 17(5):579–600, 2002.
- John AD Aston and Claudia Kirch. Detecting and estimating changes in dependent functional data. Journal of Multivariate Analysis, 109:204–220, 2012.
- John AD Aston, Davide Pigoli, and Shahin Tavakoli. Tests for separability in nonparametric covariance operators of random surfaces. The Annals of Statistics, pages 1431–1461, 2017.
- Alexander Aue, Robertas Gabrys, Lajos Horváth, and Piotr Kokoszka. Estimation of a change-point in the mean function of functional data. Journal of Multivariate Analysis, 100(10):2254–2269, 2009.
- Alexander Aue, Gregory Rice, and Ozan Sönmez. Detecting and dating structural breaks in functional data without dimension reduction. Journal of the Royal Statistical Society Series B: Statistical Methodology, 80(3):509–529, 2018.
- Alexander Aue, Gregory Rice, and Ozan Sönmez. Structural break analysis for spectrum and trace of covariance operators. Environmetrics, 31(1):e2617, 2020.
- Claudie Beaulieu, Jie Chen, and Jorge L Sarmiento. Change-point analysis as a tool to

detect abrupt climate variations. Philosophical Transactions of the Royal Society A: Mathematical, Physical and Engineering Sciences, 370(1962):1228–1249, 2012.



István Berkes, Robertas Gabrys, Lajos Horváth, and Piotr Kokoszka. Detecting changes in the mean of functional observations. Journal of the Royal Statistical Society Series B: Statistical Methodology, 71(5):927–946, 2009.

Shojaeddin Chenouri, Ahmad Mozaafari, and Gregory Rice. Robust multivariate change point analysis based on data depth. Canadian Journal of Statistics, 48(3):417–446, 2020.

Jeng-Min Chiou, Yu-Ting Chen, and Tailen Hsing. Identifying multiple changes for a functional data sequence with application to freeway traffic segmentation. The Annals of Applied Statistics, 13(3):1430–1463, 2019.

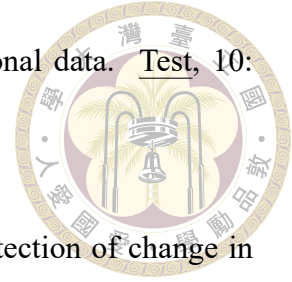
Terence Tai-Leung Chong. Estimating the locations and number of change points by the sample-splitting method. Statistical papers, 42:53–79, 2001.

Antonio Cuevas, Manuel Febrero, and Ricardo Fraiman. On the use of the bootstrap for estimating functions with functional data. Computational statistics & data analysis, 51(2):1063–1074, 2006.

Antonio Cuevas, Manuel Febrero, and Ricardo Fraiman. Robust estimation and classification for functional data via projection-based depth notions. Computational Statistics, 22(3):481–496, 2007.

Holger Dette and Kevin Kokot. Detecting relevant differences in the covariance operators of functional time series: a sup-norm approach. Annals of the Institute of Statistical Mathematics, 74(2):195–231, 2022.

Ricardo Fraiman and Graciela Muniz. Trimmed means for functional data. Test, 10: 419–440, 2001.



Oleksandr Gromenko, Piotr Kokoszka, and Matthew Reimherr. Detection of change in the spatiotemporal mean function. Journal of the Royal Statistical Society Series B: Statistical Methodology, 79(1):29–50, 2017.

Trevor Harris, Bo Li, and J Derek Tucker. Scalable multiple changepoint detection for functional data sequences. Environmetrics, 33(2):e2710, 2022.

Siegfried Hörmann and Piotr Kokoszka. Weakly dependent functional data. 2010.

Lajos Horváth and Marie Hušková. Change-point detection in panel data. Journal of Time Series Analysis, 33(4):631–648, 2012.

Huang Huang and Ying Sun. A decomposition of total variation depth for understanding functional outliers. Technometrics, 2019.

Daniela Jarušková. Testing for a change in covariance operator. Journal of Statistical Planning and Inference, 143(9):1500–1511, 2013.

Shuhao Jiao, Ron D Frostig, and Hernando Ombao. Break point detection for functional covariance. Scandinavian Journal of Statistics, 50(2):477–512, 2023.

Tze Leung Lai. Sequential changepoint detection in quality control and dynamical systems. Journal of the Royal Statistical Society: Series B (Methodological), 57(4):613–644, 1995.

Jun Li and Regina Y Liu. New nonparametric tests of multivariate locations and scales using data depth. 2004.

Regina Y Liu. On a notion of data depth based on random simplices. The Annals of Statistics, pages 405–414, 1990.

Regina Y Liu, Jesse M Parelius, and Kesar Singh. Multivariate analysis by data depth: descriptive statistics, graphics and inference,(with discussion and a rejoinder by liu and singh). The annals of statistics, 27(3):783–858, 1999.

Sara López-Pintado and Juan Romo. On the concept of depth for functional data. Journal of the American statistical Association, 104(486):718–734, 2009.

Karl Mosler. Depth statistics. Robustness and complex data structures: Festschrift in Honour of Ursula Gather, pages 17–34, 2013.

Vito MR Muggeo and Giada Adelfio. Efficient change point detection for genomic sequences of continuous measurements. Bioinformatics, 27(2):161–166, 2011.

Yue S Niu and Heping Zhang. The screening and ranking algorithm to detect dna copy number variations. The annals of applied statistics, 6(3):1306, 2012.

James O Ramsay and Bernard W Silverman. Fitting differential equations to functional data: Principal differential analysis. Springer, 2005.

Kelly Ramsay and Shoja’eddin Chenouri. Change-point detection in the covariance kernel of functional data using data depth. arXiv preprint arXiv:2112.01611, 2021.

Shawn Santo and Ping-Shou Zhong. Homogeneity tests of covariance and change-points identification for high-dimensional functional data. arXiv preprint arXiv:2005.01895, 2020.

Changliang Zou, Guanghui Wang, and Runze Li. Consistent selection of the number of change-points via sample-splitting. Annals of statistics, 48(1):413, 2020.