

Dimension reduction and manifold learning

Dimension estimation, manifold estimation

Eddie Aamari

Département de mathématiques et applications

CNRS, ENS PSL

Master IASD & MATH — Dauphine PSL

Dimension estimation

Hyperparameters

Most of the methods crucially required two types of parameters.

Bandwidths

For building neighborhood graphs

(All methods)

k -nearest neighbors

r -neighborhood

For building functions on the graph

Kernel scale σ or t

(Laplacian methods, k-PCA)

Localization radius h in local PCA

(Hessian LLE, LTSA)

Dimension

For determining the output dimension

Dimension is inherent to *dimensionality reduction*.

Intrinsic Dimension

Pioneers in intrinsic dimension estimation

This question dates back to Bennett 1969 in signal processing.

The intrinsic dimensionality of a collection of signals is defined to be equal to the number of free parameters required in a hypothetical signal generator capable of producing a close approximation to each signal in the collection. Thus defined, the dimensionality becomes a relationship between the vectors representing the signals.

Overview

See Camastra and Staiano 2016 for a recent survey.

A Statistical Remark

We are trying to estimate a **discrete quantity**

$$d \in \{0, \dots, p\}$$

If $x_1, \dots, x_n \sim_{iid} \text{Unif}_M$, we hence expect **fast estimation rates** of the form

$$\mathbb{P}(\hat{d} \neq \dim(M)) \leq C \exp(-C'n),$$

where $\hat{d} = \hat{d}(x_1, \dots, x_n)$ is some wisely chosen estimator.

A Statistical Remark

We are trying to estimate a **discrete quantity**

$$d \in \{0, \dots, p\}$$

If $x_1, \dots, x_n \sim_{iid} \text{Unif}_M$, we hence expect **fast estimation rates** of the form

$$\mathbb{P}(\hat{d} \neq \dim(M)) \leq C \exp(-C'n),$$

where $\hat{d} = \hat{d}(x_1, \dots, x_n)$ is some wisely chosen estimator.

A Take-Away Message

$$\# \{\text{Definitions of dimension}\} \asymp \# \{\text{Estimators of dimension}\} \gg 1$$

Hausdorff and Information Dimension

The **Hausdorff dimension** $\dim_H(M)$ of $M \subset \mathbb{R}^p$ is defined through

$$\Gamma_H^{(d)} := \lim_{r \rightarrow 0^+} \inf_{\substack{x_1, \dots, x_N \in S \\ r_i \leq r \\ \cup_i B(x_i, r_i) \supset M}} \sum_i r_i^d,$$

$$\text{and } \dim_H(M) := \inf \left\{ d \mid \Gamma_H^{(d)} = 0 \right\} \in [0, p]$$



Hausdorff and Information Dimension

Definition

The **information dimension** $\dim_H(P)$ of a probability measure P , is the smallest Hausdorff dimension of sets that have measure 1.

(For non-pathological cases, $\dim_H(P) = \dim_H(\text{Support}(P))$)



(Generalized) Traveling Salesman Problem

Kim, Rinaldo, and Wasserman 2019 study the testing problem

$$\mathcal{H}_0: \dim(M) = d_0 \quad \text{VS} \quad \mathcal{H}_1: \dim(M) = d_1$$

where $1 \leq d_0 < d_1 \leq p$ are fixed.

Generalized TSP Leverage of the behavior of the generalized Travelling Salesman Problem (TSP) value

$$\text{TSP}_{d_0}(\mathcal{X}) := \min_{\sigma \in \mathfrak{S}_n} \sum_{i=1}^n \|x_{\sigma(i+1)} - x_{\sigma(i)}\|^{d_0},$$

where $\mathfrak{S}_n$ is the set of permutations of $\{1, \dots, n\}$.

(Generalized) Traveling Salesman Problem

Kim, Rinaldo, and Wasserman 2019 study the testing problem

$$\mathcal{H}_0: \dim(M) = d_0 \quad \text{VS} \quad \mathcal{H}_1: \dim(M) = d_1$$

where $1 \leq d_0 < d_1 \leq p$ are fixed.

Generalized TSP Leverage of the behavior of the generalized Travelling Salesman Problem (TSP) value

$$\text{TSP}_{d_0}(\mathcal{X}) := \min_{\sigma \in \mathfrak{S}_n} \sum_{i=1}^n \|x_{\sigma(i+1)} - x_{\sigma(i)}\|^{d_0},$$

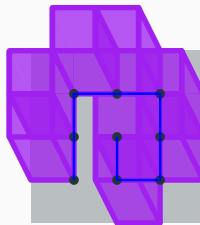
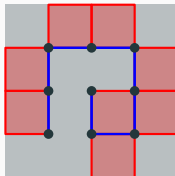
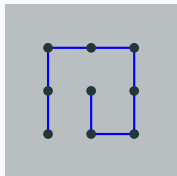
where $\mathfrak{S}_n$ is the set of permutations of $\{1, \dots, n\}$.

Intractability

Generalized TSP is NP-complete

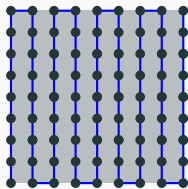
Insights Behind Generalized TSP

$$\text{TSP}_{d_0}(\mathcal{X}) := \min_{\sigma \in \mathfrak{S}_n} \sum_{i=1}^n \|x_{\sigma(i+1)} - x_{\sigma(i)}\|^{d_0}$$

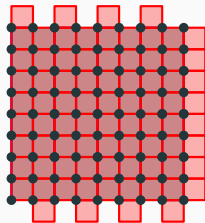


Insights Behind Generalized TSP

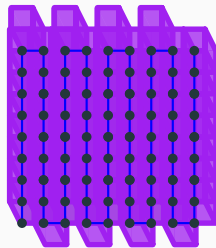
$$\text{TSP}_{d_0}(\mathcal{X}) := \min_{\sigma \in \mathfrak{S}_n} \sum_{i=1}^n \|x_{\sigma(i+1)} - x_{\sigma(i)}\|^{d_0}$$



$$d_0 < \dim(M) \\ \text{TSP}_{d_0}(\mathcal{X}) \rightarrow \infty$$



$$d_0 = \dim(M) \\ \text{TSP}_{d_0}(\mathcal{X}) = \Theta(1)$$



$$d_0 > \dim(M) \\ \text{TSP}_{d_0}(\mathcal{X}) \rightarrow 0$$

TSP Test

For $1 \leq d_0 < d_1 \leq p$ fixed,

$$\hat{d}(\mathcal{X}) := \begin{cases} d_0 & \text{if } \text{TSP}_{d_0}(\mathcal{X}) \leq C, \\ d_1 & \text{otherwise.} \end{cases}$$

Convergence Result

Theorem (Kim, Rinaldo, and Wasserman 2019)

Assume that M is $\mathcal{C}^2$ smooth and $x_1, \dots, x_n \sim_{iid} P$ uniform on M . If $\dim(M) \in \{d_0, d_1\}$, then

$$\mathbb{P}(\hat{d}(\mathcal{X}) \neq \dim(M)) \lesssim 1_{\dim(M)=d_1} \left(\frac{1}{n}\right)^{\left(\frac{d_1}{d_0}-1\right)n}$$

TSP Estimator

Define

$$\hat{d}(\mathcal{X}) := \min \{d_0 \mid \text{TSP}_{d_0}(\mathcal{X}) \leq C_{d_0}\}$$

Convergence Result

Theorem (Kim, Rinaldo, and Wasserman 2019)

If that M is $\mathcal{C}^2$ smooth and $x_1, \dots, x_n \sim_{iid} P$ uniform on M , then

$$\mathbb{P}(\hat{d}(\mathcal{X}) \neq \dim(M)) \lesssim \left(\frac{1}{n}\right)^{\frac{n}{p-1}}$$

Differential / Topological Dimension

Definition

The **topological dimension** of $M \subset \mathbb{R}^p$ is *the* dimension $\dim_R(M)$ of the model space that locally parametrizes it.

$\Rightarrow$ Local flatness

(Essentially assuming manifold structure)

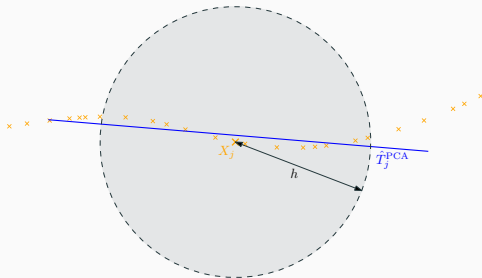
Differential / Topological Dimension

Definition

The **topological dimension** of $M \subset \mathbb{R}^p$ is *the* dimension $\dim_R(M)$ of the model space that locally parametrizes it.

$\Rightarrow$ Local flatness

(Essentially assuming manifold structure)



$\hookrightarrow$ Thresholding principal components Fukunaga and Olsen 1971

Local linear residuals

Step 1: Localize

Pick a point $x \in \mathcal{X}$, a localization radius $h > 0$, and set

$$\tilde{\mathcal{X}}_h(x) := \mathcal{X} \cap B(x, h) - x$$

Step 2: Singular Value Decomposition

Compute the **SVD** of the matrix associated with $\tilde{\mathcal{X}}_h(x)$. Store the singular values $\lambda_1 \geq \dots \geq \lambda_p$.

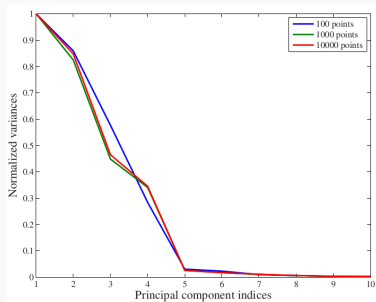
Step 3: Vary thresholds

Plot the **residual error** (or explained variance)

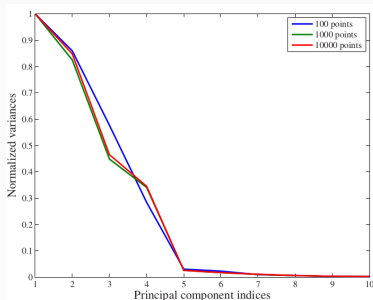
$$d \mapsto \frac{\sum_{k=d+1}^p \lambda_k}{\sum_{k=1}^p \lambda_k}$$

and search for a gap.

Illustration



Illustration



From Local to Global

In practice, need to aggregate the estimated dimensions $\hat{d}(x)$.

Trial and Error

Such a post-hoc error measurement also applies to any (local) MDS-based dimension reduction technique.

Nearest Neighbors

Instead of fixing a bandwidth, one can also regressing
k-Nearest Neighbor distances

Fukunaga and Olsen 1971, (3) show that if $x_1, \dots, x_n \sim_{iid} f(x)\lambda_d(dx)$ and $x \in \mathbb{R}^d$ with $f(x) > 0$ continuous at x , then

$$\mathbb{E}[\|x_{(k)} - x\|] \propto k^{1/d}$$

Nearest Neighbors

Instead of fixing a bandwidth, one can also regress k -Nearest Neighbor distances

Fukunaga and Olsen 1971, (3) show that if $x_1, \dots, x_n \sim_{iid} f(x)\lambda_d(dx)$ and $x \in \mathbb{R}^d$ with $f(x) > 0$ continuous at x , then

$$\mathbb{E}[\|x_{(k)} - x\|] \propto k^{1/d}$$

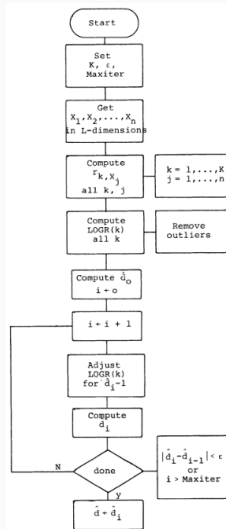


Fig. 1. Procedure for computing $\hat{d}$.

Correlation Dimension

Leveraging local neighborhood properties can also be done by noticing that if $x, x' \stackrel{\mathbb{P}}{\sim} \lambda_d$,

$$\mathbb{P}(\|x - x'\| \leq r) \asymp r^d.$$

This leads to the **correlation dimension**, based on

$$\text{Cor}_r^{(2)}(P) := \mathbb{P}_{x, x' \stackrel{\mathbb{P}}{\sim} P}(\|x - x'\| \leq r),$$

and defined as

$$\text{dim}_{\text{cor},2}(P) := \lim_{r \rightarrow 0} \frac{\log \text{Cor}_r^{(2)}(P)}{\log r}$$

Correlation Dimension

Estimator

Given $x_1, \dots, x_n \sim_{iid} P$, consider the U-statistic

$$\widehat{\text{Cor}}_r^{(2)} := \frac{2}{n(n-1)} \sum_{i < j} \mathbf{1}_{\|x_i - x_j\| \leq r},$$

with associate dimension estimator

$$\hat{d}_{\text{cor},2} := \lim_{r \rightarrow 0} \frac{\log \widehat{\text{Cor}}_r^{(2)}}{\log r}$$

Convergence

See results in Kégl [2002](#) and Arias-Castro, Chen, and Lerman [2011](#), Section 2.2.1.

Covering Number

Definition

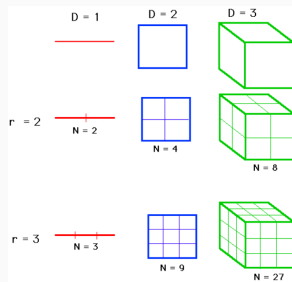
Given $M \subset \mathbb{R}^p$, the r -covering number of M is

$$\text{cv}_M(r) := \min \left\{ N \mid \exists z_1, \dots, z_N \in \mathbb{R}^p \text{ s.t. } M \subset \bigcup_{i=1}^N B(x_i, r) \right\}$$

The r -dimension of M is $\dim_r(M) := \frac{\log \text{cv}_M(r)}{-\log r}$.

For $M = [0, 1]^d$,

$$\text{cv}_M(r) \asymp \left(\frac{1}{r} \right)^d$$



Definition

The **Minkowski (or Capacity) dimension** of M is

$$\dim_{\text{Min}}(M) := \limsup_{r \rightarrow 0} \dim_r(M).$$

Insights

If $\dim_{\text{Min}}(M) = d$, we expect that

$$\log \text{cv}_M(r) \sim_{r \rightarrow 0} -d \log r$$

Regression

We can regress

$$r \mapsto \log \text{cv}_{\mathcal{X}}(r)$$

Two-Scales Estimation

Instead of regression, Kégl 2002 uses the fact that for all small $r_1 < r_2$,

$$\frac{\log \text{cv}_M(r_1) - \log \text{cv}_M(r_2)}{\log r_2 - \log r_1} \simeq \frac{-d \log r_1 - (-d \log r_2)}{\log r_2 - \log r_1} = d.$$

Limitations

Still a choice of bandwidth(s) parameter(s).

Costly to compute directly on data (involves covering numbers).

Wayaround

Try to observe the dimension indirectly on a simpler object.

Partial Coverings

Weed and Bach 2019 introduce the Wasserstein dimension.

Idea

When working with measures instead of sets, it is convenient to be able to ignore a small fraction of the mass.

Definition

The (r, τ) -covering number of a probability measure P on $\mathbb{R}^p$ is

$$\text{cv}_P(r, \tau) := \min \{ \text{cv}_S(r) \mid P(S) \geq 1 - \tau \}.$$

Its (r, τ) -dimension is

$$\dim_{r, \tau}(P) := \frac{\log \text{cv}_P(r, \tau)}{-\log r}.$$

Definition (Weed and Bach 2019)

The upper and lower Wasserstein dimensions of P are respectively

$$\begin{aligned}\bar{d}^{(p)}(P) &:= \inf \left\{ s > 2p \mid \limsup_{r \rightarrow 0} \dim_{r, r^{\frac{sp}{s-2p}}}(P) \leq s \right\} \\ \underline{d}^{(p)}(P) &:= \lim_{\tau \rightarrow 0} \liminf_{r \rightarrow 0} \dim_{r, \tau}(P).\end{aligned}$$

Definition (Weed and Bach 2019)

The **upper and lower Wasserstein dimensions** of P are respectively

$$\begin{aligned}\bar{d}^{(p)}(P) &:= \inf \left\{ s > 2p \mid \limsup_{r \rightarrow 0} \dim_{r, r^{\frac{sp}{s-2p}}}(P) \leq s \right\} \\ \underline{d}^{(p)}(P) &:= \lim_{\tau \rightarrow 0} \liminf_{r \rightarrow 0} \dim_{r, \tau}(P).\end{aligned}$$

Links with other dimensions

If $P(B(x, r)) \asymp r^d$, then $\bar{d}^{(p)}(P) = d = \underline{d}^{(p)}(P)$

Generalizable to arbitrary ambient metric space.

Convergence of the Empirical Distribution

Theorem (Weed and Bach 2019)

Let $p \geq 1$. Assume that $x_1, \dots, x_n \sim_{iid} P$ on $\mathbb{R}^d$. and write $P_n := n^{-1} \sum_{i=1}^n \delta_{x_i}$ for the empirical measure.

If $s > \overline{\dim}^{(p)}(P)$, then $\mathbb{E} [W_p(P, P_n)] \lesssim n^{-1/s}$.

If $t < \underline{\dim}^{(p)}(P)$, then $\mathbb{E} [W_p(P, P_n)] \gtrsim n^{-1/t}$.

Convergence of the Empirical Distribution

Theorem (Weed and Bach 2019)

Let $p \geq 1$. Assume that $x_1, \dots, x_n \sim_{iid} P$ on $\mathbb{R}^d$. and write $P_n := n^{-1} \sum_{i=1}^n \delta_{x_i}$ for the empirical measure.

If $s > \overline{\dim}^{(p)}(P)$, then $\mathbb{E}[W_p(P, P_n)] \lesssim n^{-1/s}$.

If $t < \underline{\dim}^{(p)}(P)$, then $\mathbb{E}[W_p(P, P_n)] \gtrsim n^{-1/t}$.

The upper bound comes from a spatial **dyadic decomposition**.

The lower bound holds for all distribution P_n supported on n Dirac. It arises from a **quantization** argument.

Fine results on $W_p(P, P_n)$ in Dedecker and Merlevède 2019.

Block et al. 2021 leverage this sharp behavior as follows.

Bootstrap-Style Method

Given $0 < \alpha < 1$, subsample:

(need $2(1 + \alpha)n$ independent sample)

P_n, P'_n each arising from n observations each

$P_{\alpha n}, P'_{\alpha n}$ each arising from $\alpha n < n$ observations each

As

$$W_1(P_{[\alpha]n}, P'_{[\alpha]n}) \asymp W_1(P, P_{[\alpha]n}) \asymp (1/([\alpha]n))^{1/d},$$

take

$$\hat{d}_W := \frac{\log \alpha}{\log W_1(P_n, P'_n) - \log W_1(P_{\alpha n}, P'_{\alpha n})}$$

Which Wasserstein Metric?

Possibility to use the (estimated) geodesic metric in Wasserstein.

Geometric inference

Take a step back

Throughout, we have tried to **embed** points $\mathcal{X} \subset \mathbb{R}^p$ to $\mathcal{Y} \subset \mathbb{R}^d$ while **preserving the geometry of $\mathcal{X}$** .

If we assume that $\mathcal{X} \subset M$ are sample from a **submanifold** $M \subset \mathbb{R}^p$:

Preserving the geometry of $\mathcal{X}$

$\Leftrightarrow$

$$d_M(x_i, x_j) \simeq \|y_i - y_j\|$$

Take a step back

Throughout, we have tried to **embed** points $\mathcal{X} \subset \mathbb{R}^p$ to $\mathcal{Y} \subset \mathbb{R}^d$ while **preserving the geometry of $\mathcal{X}$** .

If we assume that $\mathcal{X} \subset M$ are sample from a **submanifold** $M \subset \mathbb{R}^p$:

Preserving the geometry of $\mathcal{X}$

$\Leftrightarrow$

$$d_M(x_i, x_j) \simeq \|y_i - y_j\|$$

The **geodesic distance** on M

(or **shortest-path distance**)

$$\begin{aligned} d_M: M \times M &\longrightarrow \mathbb{R}_{\geq 0} \cup \{\infty\} \\ (x, y) &\longmapsto \inf_{\substack{\gamma_{x \rightarrow y} \subset M \\ \mathcal{C}^1 \text{ curve}}} \int \|\gamma'_{x \rightarrow y}\| \end{aligned}$$

Take a step back

Throughout, we have tried to **embed** points $\mathcal{X} \subset \mathbb{R}^p$ to $\mathcal{Y} \subset \mathbb{R}^d$ while **preserving the geometry of $\mathcal{X}$** .

If we assume that $\mathcal{X} \subset M$ are sample from a **submanifold** $M \subset \mathbb{R}^p$:

Preserving the geometry of $\mathcal{X}$

$\Leftrightarrow$

$$d_M(x_i, x_j) \simeq \|y_i - y_j\|$$

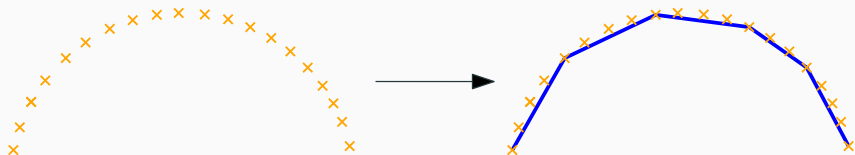
The **geodesic distance** on M

(or **shortest-path distance**)

$$\begin{aligned} d_M: M \times M &\longrightarrow \mathbb{R}_{\geq 0} \cup \{\infty\} \\ (x, y) &\longmapsto \inf_{\substack{\gamma_{x \rightarrow y} \subset M \\ \mathcal{C}^1 \text{ curve}}} \int \|\gamma'_{x \rightarrow y}\| \end{aligned}$$

What about only estimating d_M without embedding?

Manifold Estimation



Given

- a point cloud $\mathcal{X} = \{x_1, \dots, x_n\}$ close to $M \subset \mathbb{R}^p$
- with M of intrinsic dimension d

find a 'suitable' $\hat{M} \subset \mathbb{R}^p$ of dimension d such that

$\hat{M}$ is 'close' to M .

Manifold Estimation



Given

- a point cloud $\mathcal{X} = \{x_1, \dots, x_n\}$ close to $M \subset \mathbb{R}^p$
- with M of intrinsic dimension d

find a 'suitable' $\hat{M} \subset \mathbb{R}^p$ of dimension d such that

$\hat{M}$ is 'close' to M .

Does estimating M help to do better dimension reduction?

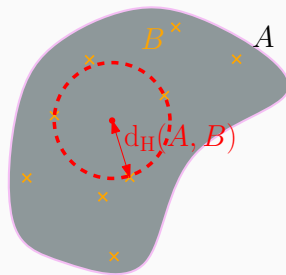
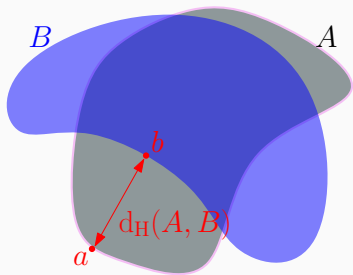
Hausdorff Distance

Definition (Hausdorff Distance)

The *Hausdorff distance* between two compact sets $A, B \subset \mathbb{R}^p$ is

$$d_H(A, B) = \|d(\cdot, A) - d(\cdot, B)\|_\infty,$$

where $d(x, C) := \inf_{c \in C} \|x - c\|$ is the distance to $C \subset \mathbb{R}^p$.

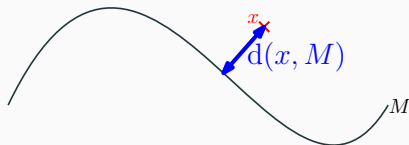


Disambiguation

- The *distance function* to M :

(used to identify sets as functions)

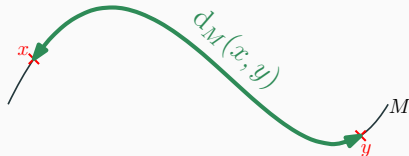
$$\begin{aligned} d(\cdot, M): \mathbb{R}^p &\longrightarrow \mathbb{R}_{\geq 0} \\ x &\longmapsto \min_{p \in M} \|x - p\| \end{aligned}$$



- The *geodesic distance* on M :

(used to characterize the geometry of sets)

$$\begin{aligned} d_M: M \times M &\longrightarrow \mathbb{R}_{\geq 0} \cup \{\infty\} \\ (x, y) &\longmapsto \inf_{\substack{\gamma_{x \rightarrow y} \subset M \\ \mathcal{C}^1 \text{ curve}}} \int \|\gamma'_{x \rightarrow y}\| \end{aligned}$$



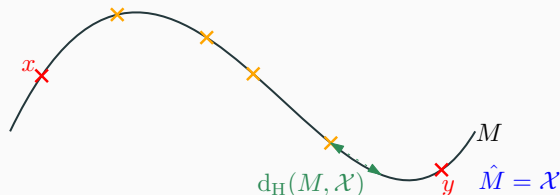
From Manifold Estimation to Metric Learning

Theorem (Aamari, Berenfeld, Levrard – 2023)

Assume that $M \subset \mathbb{R}^p$ is $\mathcal{C}^2$ -smooth. Then there exists $\text{rch}_M > 0$ such that for all $\hat{M} \subset \mathbb{R}^p$ such that $d_H(M, \hat{M}) \leq \varepsilon < \text{rch}_M/2$,

$$\sup_{x \neq y \in M} |d_M(x, y) - d_{(\hat{M})^\varepsilon}(x, y)| \lesssim \varepsilon,$$

where $(\hat{M})^\varepsilon := \{u \in \mathbb{R}^p \mid d(u, \hat{M}) \leq \varepsilon\}$ is the ε -offset of M .



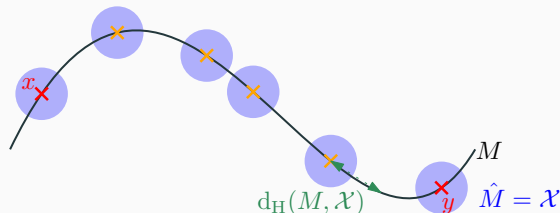
From Manifold Estimation to Metric Learning

Theorem (Aamari, Berenfeld, Levrard – 2023)

Assume that $M \subset \mathbb{R}^p$ is $\mathcal{C}^2$ -smooth. Then there exists $\text{rch}_M > 0$ such that for all $\hat{M} \subset \mathbb{R}^p$ such that $d_H(M, \hat{M}) \leq \varepsilon < \text{rch}_M/2$,

$$\sup_{x \neq y \in M} |d_M(x, y) - d_{(\hat{M})^\varepsilon}(x, y)| \lesssim \varepsilon,$$

where $(\hat{M})^\varepsilon := \{u \in \mathbb{R}^p \mid d(u, \hat{M}) \leq \varepsilon\}$ is the ε -offset of M .



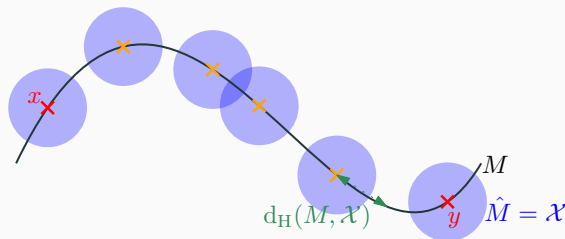
From Manifold Estimation to Metric Learning

Theorem (Aamari, Berenfeld, Levrard – 2023)

Assume that $M \subset \mathbb{R}^p$ is $\mathcal{C}^2$ -smooth. Then there exists $\text{rch}_M > 0$ such that for all $\hat{M} \subset \mathbb{R}^p$ such that $d_H(M, \hat{M}) \leq \varepsilon < \text{rch}_M/2$,

$$\sup_{x \neq y \in M} |d_M(x, y) - d_{(\hat{M})^\varepsilon}(x, y)| \lesssim \varepsilon,$$

where $(\hat{M})^\varepsilon := \{u \in \mathbb{R}^p \mid d(u, \hat{M}) \leq \varepsilon\}$ is the ε -offset of M .



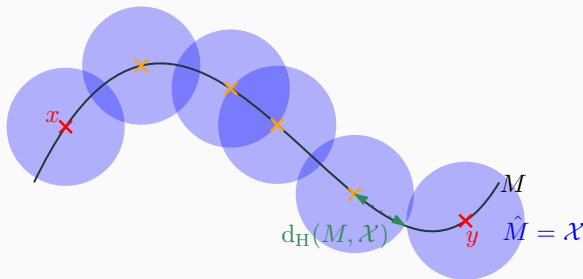
From Manifold Estimation to Metric Learning

Theorem (Aamari, Berenfeld, Levrard – 2023)

Assume that $M \subset \mathbb{R}^p$ is $\mathcal{C}^2$ -smooth. Then there exists $\text{rch}_M > 0$ such that for all $\hat{M} \subset \mathbb{R}^p$ such that $d_H(M, \hat{M}) \leq \varepsilon < \text{rch}_M/2$,

$$\sup_{x \neq y \in M} |d_M(x, y) - d_{(\hat{M})^\varepsilon}(x, y)| \lesssim \varepsilon,$$

where $(\hat{M})^\varepsilon := \{u \in \mathbb{R}^p \mid d(u, \hat{M}) \leq \varepsilon\}$ is the ε -offset of M .



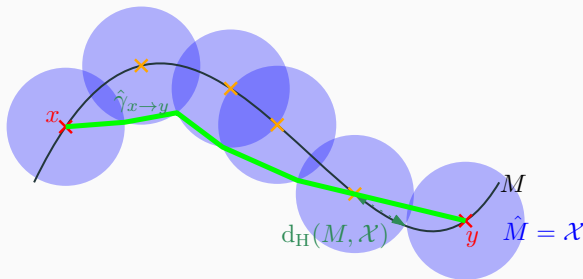
From Manifold Estimation to Metric Learning

Theorem (Aamari, Berenfeld, Levrard – 2023)

Assume that $M \subset \mathbb{R}^p$ is $\mathcal{C}^2$ -smooth. Then there exists $\text{rch}_M > 0$ such that for all $\hat{M} \subset \mathbb{R}^p$ such that $d_H(M, \hat{M}) \leq \varepsilon < \text{rch}_M/2$,

$$\sup_{x \neq y \in M} |d_M(x, y) - d_{(\hat{M})^\varepsilon}(x, y)| \lesssim \varepsilon,$$

where $(\hat{M})^\varepsilon := \{u \in \mathbb{R}^p \mid d(u, \hat{M}) \leq \varepsilon\}$ is the ε -offset of M .



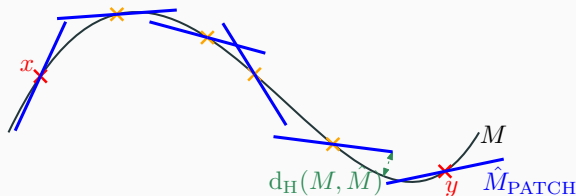
From Manifold Estimation to Metric Learning

Theorem (Aamari, Berenfeld, Levrard – 2023)

Assume that $M \subset \mathbb{R}^p$ is $\mathcal{C}^2$ -smooth. Then there exists $\text{rch}_M > 0$ such that for all $\hat{M} \subset \mathbb{R}^p$ such that $d_H(M, \hat{M}) \leq \varepsilon < \text{rch}_M/2$,

$$\sup_{x \neq y \in M} |d_M(x, y) - d_{(\hat{M})^\varepsilon}(x, y)| \lesssim \varepsilon,$$

where $(\hat{M})^\varepsilon := \{u \in \mathbb{R}^p \mid d(u, \hat{M}) \leq \varepsilon\}$ is the ε -offset of M .



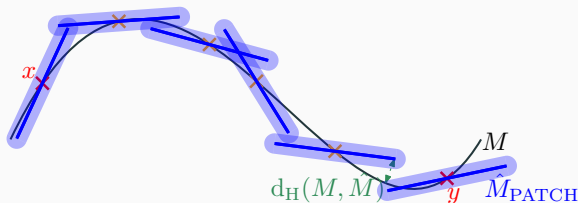
From Manifold Estimation to Metric Learning

Theorem (Aamari, Berenfeld, Levrard – 2023)

Assume that $M \subset \mathbb{R}^p$ is $\mathcal{C}^2$ -smooth. Then there exists $\text{rch}_M > 0$ such that for all $\hat{M} \subset \mathbb{R}^p$ such that $d_H(M, \hat{M}) \leq \varepsilon < \text{rch}_M/2$,

$$\sup_{x \neq y \in M} |d_M(x, y) - d_{(\hat{M})^\varepsilon}(x, y)| \lesssim \varepsilon,$$

where $(\hat{M})^\varepsilon := \{u \in \mathbb{R}^p \mid d(u, \hat{M}) \leq \varepsilon\}$ is the ε -offset of M .



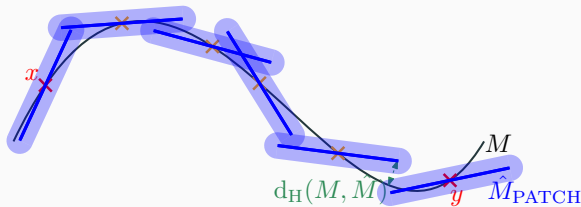
From Manifold Estimation to Metric Learning

Theorem (Aamari, Berenfeld, Levrard – 2023)

Assume that $M \subset \mathbb{R}^p$ is $\mathcal{C}^2$ -smooth. Then there exists $\text{rch}_M > 0$ such that for all $\hat{M} \subset \mathbb{R}^p$ such that $d_H(M, \hat{M}) \leq \varepsilon < \text{rch}_M/2$,

$$\sup_{x \neq y \in M} |d_M(x, y) - d_{(\hat{M})^\varepsilon}(x, y)| \lesssim \varepsilon,$$

where $(\hat{M})^\varepsilon := \{u \in \mathbb{R}^p \mid d(u, \hat{M}) \leq \varepsilon\}$ is the ε -offset of M .



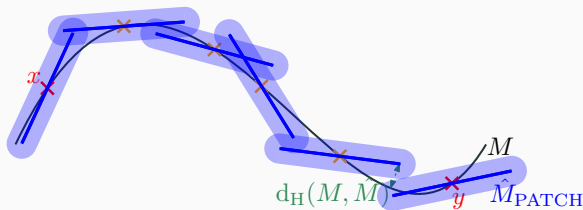
From Manifold Estimation to Metric Learning

Theorem (Aamari, Berenfeld, Levrard – 2023)

Assume that $M \subset \mathbb{R}^p$ is $\mathcal{C}^2$ -smooth. Then there exists $\text{rch}_M > 0$ such that for all $\hat{M} \subset \mathbb{R}^p$ such that $d_H(M, \hat{M}) \leq \varepsilon < \text{rch}_M/2$,

$$\sup_{x \neq y \in M} |d_M(x, y) - d_{(\hat{M})^\varepsilon}(x, y)| \lesssim \varepsilon,$$

where $(\hat{M})^\varepsilon := \{u \in \mathbb{R}^p \mid d(u, \hat{M}) \leq \varepsilon\}$ is the ε -offset of M .



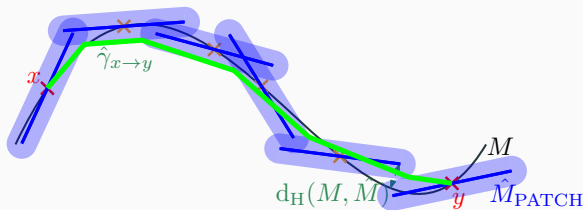
From Manifold Estimation to Metric Learning

Theorem (Aamari, Berenfeld, Levrard – 2023)

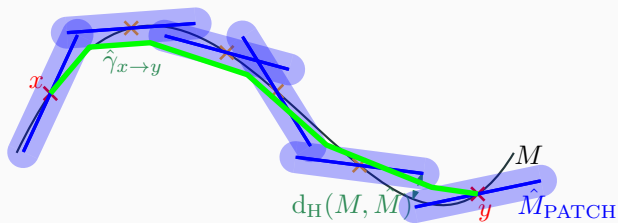
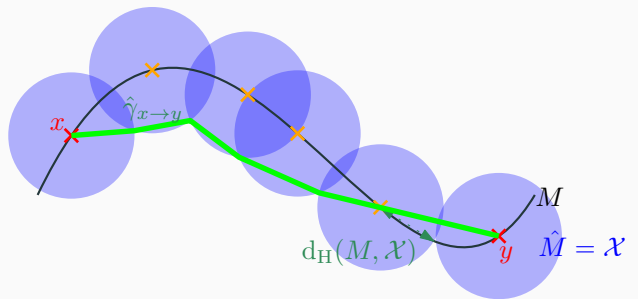
Assume that $M \subset \mathbb{R}^p$ is $\mathcal{C}^2$ -smooth. Then there exists $\text{rch}_M > 0$ such that for all $\hat{M} \subset \mathbb{R}^p$ such that $d_H(M, \hat{M}) \leq \varepsilon < \text{rch}_M/2$,

$$\sup_{x \neq y \in M} |d_M(x, y) - d_{(\hat{M})^\varepsilon}(x, y)| \lesssim \varepsilon,$$

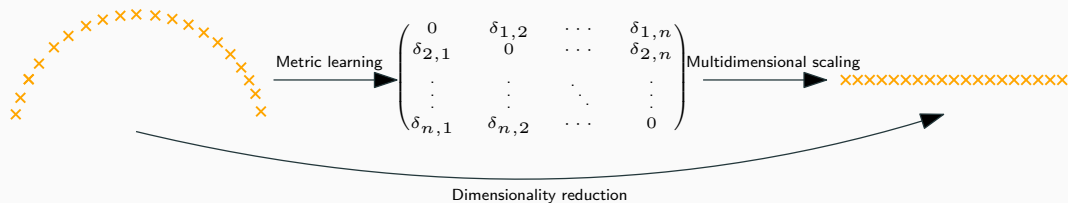
where $(\hat{M})^\varepsilon := \{u \in \mathbb{R}^p \mid d(u, \hat{M}) \leq \varepsilon\}$ is the ε -offset of M .



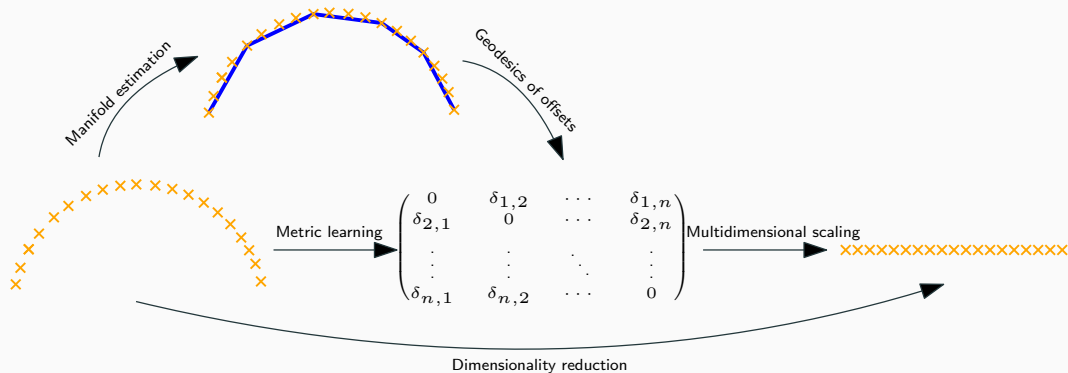
Better Manifold Estimation, Better Metric Learning



Manifold estimation $\Rightarrow$ Dimensionality reduction

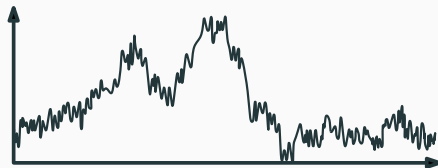
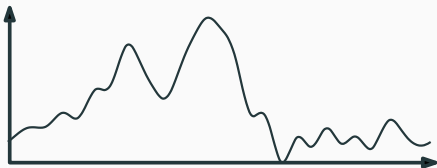


Manifold estimation $\Rightarrow$ Dimensionality reduction



Regularity in nonparametric geometric problems

Regularity

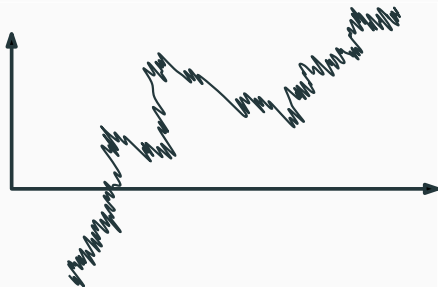
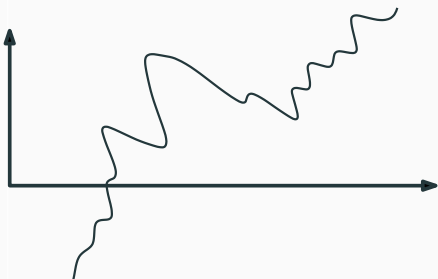


Usual regularity classes (Hölder, Sobolev, Besov) control increments

$$\|f(x) - f(y)\| \leq L \|x - y\|^\beta.$$

(L, β) drives the difficulty of the statistical problem.

Regularity Without Coordinates?



Usual regularity classes (Hölder, Sobolev, Besov) control increments

$$\|f(x) - f(y)\| \leq L \|x - y\|^\beta.$$

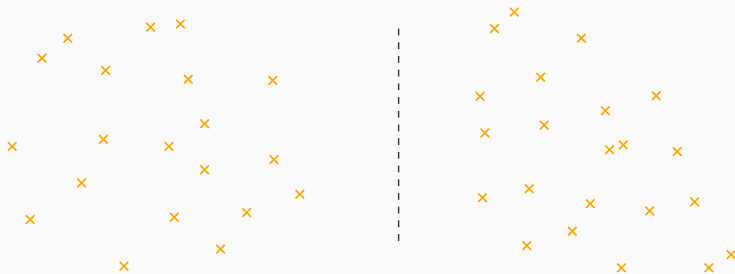
(L, β) drives the difficulty of the statistical problem.

Without natural coordinates, “ $\|f(x) - f(y)\|$ ” = ?

Support Estimation

Data: A n -sample $X_1, \dots, X_n \stackrel{i.i.d.}{\sim} P$.

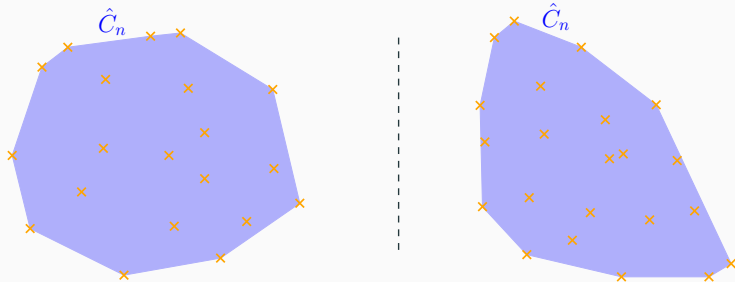
Goal: Estimate the set $C = \text{Support}(P) = \bigcap_{\substack{K \subset \mathbb{R}^p \text{ closed} \\ P(K)=1}} K$.



Support Estimation

Data: A n -sample $X_1, \dots, X_n \stackrel{i.i.d.}{\sim} P$.

Goal: Estimate the set $C = \text{Support}(P) = \bigcap_{\substack{K \subset \mathbb{R}^p \text{ closed} \\ P(K)=1}} K$.



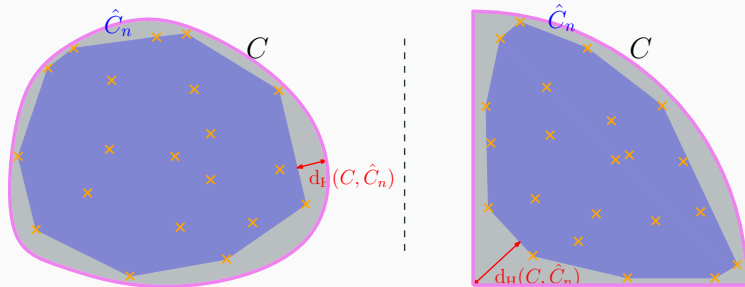
If we know (by advance) that C is convex, a good candidate is

$$\hat{C}_n := \text{Conv}(\{X_1, \dots, X_n\}).$$

Support Estimation

Data: A n -sample $X_1, \dots, X_n \stackrel{i.i.d.}{\sim} P$.

Goal: Estimate the set $C = \text{Support}(P) = \bigcap_{\substack{K \subset \mathbb{R}^p \text{ closed} \\ P(K)=1}} K$.



If we know (by advance) that C is convex, a good candidate is

$$\hat{C}_n := \text{Conv}(\{X_1, \dots, X_n\}).$$

Support Estimation: Convex Case(s)

Theorem (Dümbgen, Walther – 1996)

Assume that $P = \text{Unif}_C$ is uniform over the convex set $C \subset \mathbb{R}^p$. Write

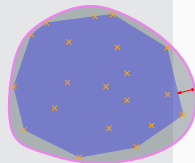
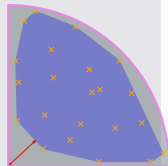
$$\mathbb{X}_n := \{X_1, \dots, X_n\}, \quad \text{and} \quad \hat{C}_n = \text{Conv}(\mathbb{X}_n).$$

– Then,

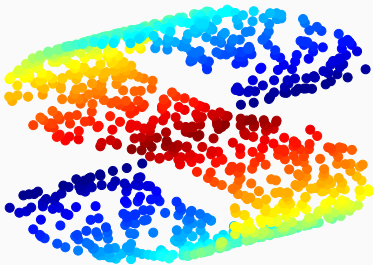
$$d_H(C, \mathbb{X}_n) \asymp d_H(C, \hat{C}_n) = O\left(\frac{\log n}{n}\right)^{\frac{1}{D}} \text{ a.s.}$$

– If in addition ∂C is $\mathcal{C}^2$,

$$d_H(C, \hat{C}_n) = O\left(\frac{\log n}{n}\right)^{\frac{2}{D+1}} \text{ a.s.}$$

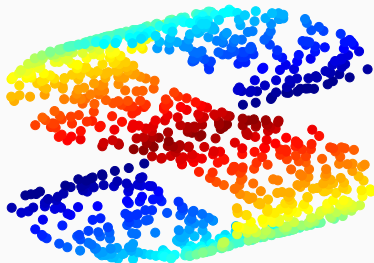


Beyond Convexity



How to model the support of these data?

- Low-dimensional and curved $\rightarrow$ Submanifold of $\mathbb{R}^p$.
- Not convex, but *locally* around it the projection uniquely defined.



How to model the support of these data?

- Low-dimensional and curved $\rightarrow$ Submanifold of $\mathbb{R}^p$.
- Not convex, but *locally* around it the projection uniquely defined.

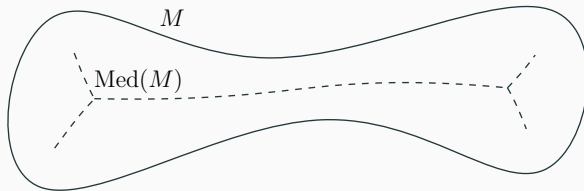
Reminder: For a closed set $C \subset \mathbb{R}^p$,

$C \subset \mathbb{R}^p$ is convex $\Leftrightarrow$ Every $z \in \mathbb{R}^p$ has a unique nearest neighbor on C
i.e. $\exists!$ $\pi_C(z) \in C$ with $\|z - \pi_C(z)\| = d(z, C)$.

Medial Axis

The **medial axis** of $M \subset \mathbb{R}^p$ is the set of points that have ≥ 2 nearest neighbors on M :

$$\text{Med}(M) := \{z \in \mathbb{R}^p \mid z \text{ has several nearest neighbors on } M\}.$$



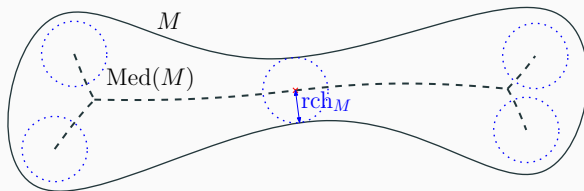
Medial axis of a curve

Reach

For a closed subset $M \subset \mathbb{R}^p$, the **reach** rch_M of M is the least distance to its medial axis:

$$\text{rch}_M := \inf_{x \in M} d(x, \text{Med}(M)),$$

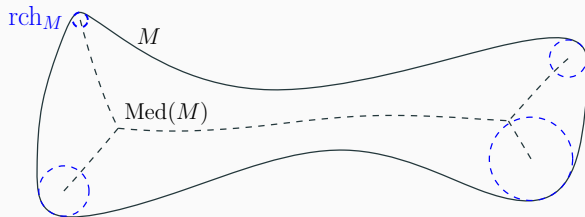
where for all $x \in \mathbb{R}^p$, $d(x, K) := \inf_{p \in K} \|x - p\|$.



One can also flip the formula:

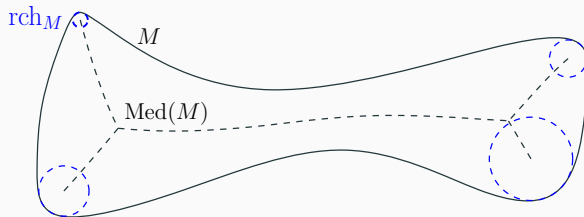
$$\text{rch}_M = \inf_{z \in \text{Med}(M)} d(z, M).$$

Local Regularity



High curvature $\Leftrightarrow$ Small radius of curvature $\Rightarrow \text{rch}_M \ll 1$.

Local Regularity



High curvature $\Leftrightarrow$ Small radius of curvature $\Rightarrow \text{rch}_M \ll 1$.

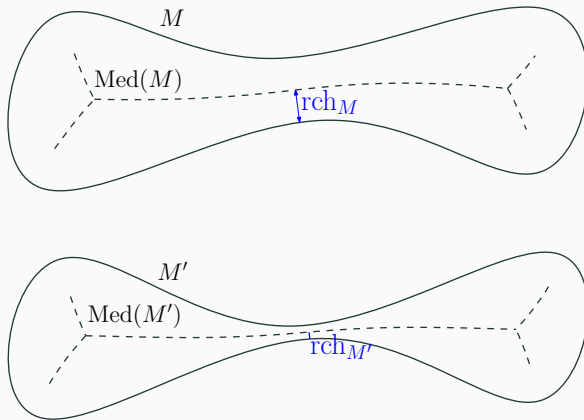
Proposition (Federer – 1959, Niyogi et al. – 2006)

Let II_x^M denote the second fundamental form of M . For all unit tangent vector $v \in T_x M$,

$$\|II_x^M(v, v)\| \leq 1/\text{rch}_M.$$

As a consequence, the sectional curvatures κ of M satisfy

$$-2/\text{rch}_M^2 \leq \kappa \leq 1/\text{rch}_M^2.$$



Narrow bottleneck structure $\Rightarrow \text{rch}_M \ll 1$.

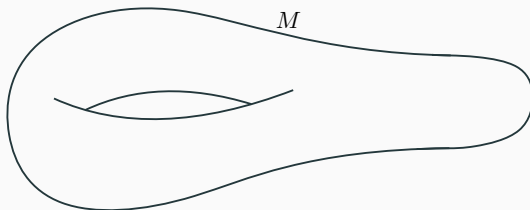
Noiseless manifold estimation

Boundariless Statistical Model

$X_1, \dots, X_n \stackrel{i.i.d.}{\sim} P$, where $M = \text{Support}(P) \subset \mathbb{R}^p$ satisfies:

- M is a compact connected d -dimensional submanifold,
- M has **no boundary**,
- $\text{rch}_M \geq \text{rch}_{\min} > 0$,
- P is (almost) the uniform distribution on M .

The set of distributions satisfying these conditions is denoted by $\mathcal{P}$.



A Reconstruction Theorem

Theorem (Aamari and Levrard 2019)

If $P \in \mathcal{P}$, one can compute an estimator $\hat{M}_{\text{PATCH}}$ based on data points $\mathbb{X}_n$ such that w.h.p.,

$$d_H(M, \hat{M}_{\text{PATCH}}) \leq C \left(\frac{\log n}{n} \right)^{2/d}.$$

Here, $C = C_{\text{rch}_{\min}, d}$ does not depend on the ambient dimension $p \gg d$.

A Reconstruction Theorem

Theorem (Aamari and Levrard 2019)

If $P \in \mathcal{P}$, one can compute an estimator $\hat{M}_{\text{PATCH}}$ based on data points $\mathbb{X}_n$ such that w.h.p.,

$$d_H(M, \hat{M}_{\text{PATCH}}) \leq C \left(\frac{\log n}{n} \right)^{2/d}.$$

Here, $C = C_{\text{rch}_{\min}, d}$ does not depend on the ambient dimension $p \gg d$.

→ Other estimators achieving the same Hausdorff rate:

- Empirical risk manifold minimizer Genovese et al. 2012b
- Local Tangent Delaunay triangulation Aamari and Levrard 2018
- Local convex hulls Divol 2021

Ingredient I: Approximation Theory

Theorem (Aamari and Levrard 2019)

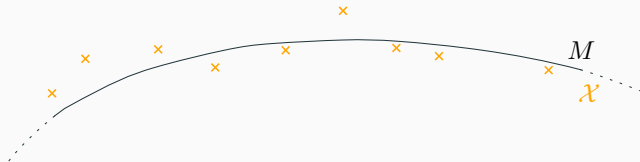
For $\Delta \lesssim \text{rch}_{\min}$, assume that we have a point cloud $\mathcal{X} \subset \mathbb{R}^p$ that is:

- *close to M* : $\max_{x \in \mathcal{X}} d(x, M) \lesssim \Delta^2 / \text{rch}_{\min},$
- *a covering of M* : $\sup_{p \in M} d(p, \mathcal{X}) \lesssim \Delta,$

together with a family $\mathbb{T}_{\mathcal{X}}$ of linear spaces that

- *approximate tangent spaces*: $\max_{x \in \mathcal{X}} \angle(T_{\pi_M(x)} M, T_x) \lesssim \Delta / \text{rch}_{\min}$

One can build a local linear estimator $\hat{M}_{\text{PATCH}}(\mathcal{X}, \mathbb{T}_{\mathcal{X}})$ such that $d_H(M, \hat{M}_{\text{PATCH}}) \lesssim \Delta^2 / \text{rch}_{\min}.$



Ingredient I: Approximation Theory

Theorem (Aamari and Levrard 2019)

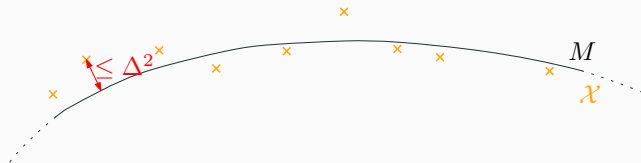
For $\Delta \lesssim \text{rch}_{\min}$, assume that we have a point cloud $\mathcal{X} \subset \mathbb{R}^p$ that is:

- *close to M* : $\max_{x \in \mathcal{X}} d(x, M) \lesssim \Delta^2 / \text{rch}_{\min},$
- *a covering of M* : $\sup_{p \in M} d(p, \mathcal{X}) \lesssim \Delta,$

together with a family $\mathbb{T}_{\mathcal{X}}$ of linear spaces that

- *approximate tangent spaces*: $\max_{x \in \mathcal{X}} \angle(T_{\pi_M(x)} M, T_x) \lesssim \Delta / \text{rch}_{\min}$

One can build a local linear estimator $\hat{M}_{\text{PATCH}}(\mathcal{X}, \mathbb{T}_{\mathcal{X}})$ such that $d_H(M, \hat{M}_{\text{PATCH}}) \lesssim \Delta^2 / \text{rch}_{\min}.$



Ingredient I: Approximation Theory

Theorem (Aamari and Levrard 2019)

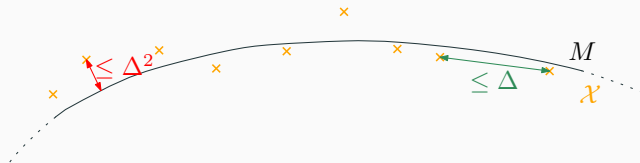
For $\Delta \lesssim \text{rch}_{\min}$, assume that we have a point cloud $\mathcal{X} \subset \mathbb{R}^p$ that is:

- *close to M* : $\max_{x \in \mathcal{X}} d(x, M) \lesssim \Delta^2 / \text{rch}_{\min},$
- *a covering of M* : $\sup_{p \in M} d(p, \mathcal{X}) \lesssim \Delta,$

together with a family $\mathbb{T}_{\mathcal{X}}$ of linear spaces that

- *approximate tangent spaces*: $\max_{x \in \mathcal{X}} \angle(T_{\pi_M(x)} M, T_x) \lesssim \Delta / \text{rch}_{\min}$

One can build a local linear estimator $\hat{M}_{\text{PATCH}}(\mathcal{X}, \mathbb{T}_{\mathcal{X}})$ such that $d_H(M, \hat{M}_{\text{PATCH}}) \lesssim \Delta^2 / \text{rch}_{\min}.$



Ingredient I: Approximation Theory

Theorem (Aamari and Levrard 2019)

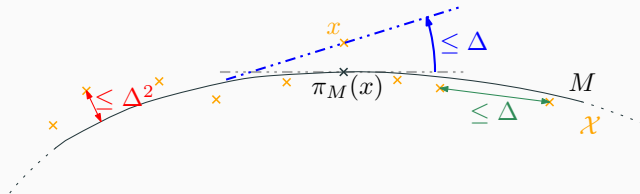
For $\Delta \lesssim \text{rch}_{\min}$, assume that we have a point cloud $\mathcal{X} \subset \mathbb{R}^p$ that is:

- *close to M* : $\max_{x \in \mathcal{X}} d(x, M) \lesssim \Delta^2 / \text{rch}_{\min},$
- *a covering of M* : $\sup_{p \in M} d(p, \mathcal{X}) \lesssim \Delta,$

together with a family $\mathbb{T}_{\mathcal{X}}$ of linear spaces that

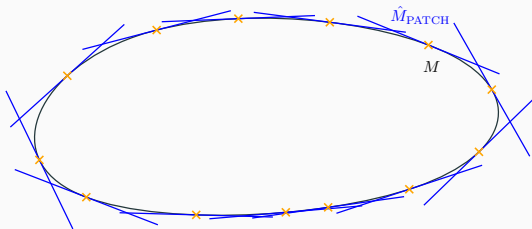
- *approximate tangent spaces*: $\max_{x \in \mathcal{X}} \angle(T_{\pi_M(x)} M, T_x) \lesssim \Delta / \text{rch}_{\min}$

One can build a local linear estimator $\hat{M}_{\text{PATCH}}(\mathcal{X}, \mathbb{T}_{\mathcal{X}})$ such that $d_H(M, \hat{M}_{\text{PATCH}}) \lesssim \Delta^2 / \text{rch}_{\min}.$



Ingredient I: Approximation Theory

$$\hat{M}_{\text{PATCH}} := \bigcup_{x \in \mathcal{X}} B_{T_x}(0, \Delta).$$



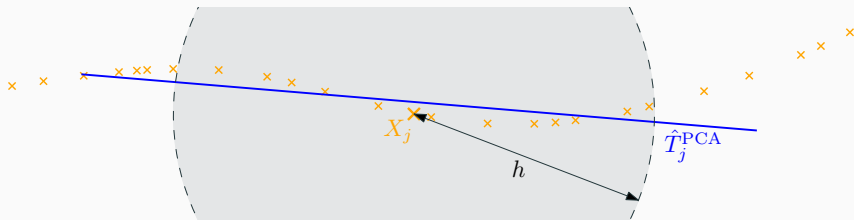
Ingredient II: Local PCA

Define $\hat{T}_j^{\text{PCA}}$ to be a minimizer of

$$\hat{T}_j^{\text{PCA}} \in \arg \min_T P_n^{(j)} \left[\|x - \pi_T(x)\|^2 \mathbf{1}_{B(0,h)}(x) \right],$$

where:

- $P_n^{(j)}$ denotes the integration with respect to $\frac{1}{n} \sum_{\ell \neq j} \delta_{X_\ell - X_j}$,
- T ranges in the set of d -planes of $\mathbb{R}^p$.



Ingredient II: Local PCA

Define $\hat{T}_j^{\text{PCA}}$ to be a minimizer of

$$\hat{T}_j^{\text{PCA}} \in \arg \min_T P_n^{(j)} \left[\|x - \pi_T(x)\|^2 \mathbf{1}_{B(0,h)}(x) \right],$$

where:

- $P_n^{(j)}$ denotes the integration with respect to $\frac{1}{n} \sum_{\ell \neq j} \delta_{X_\ell - X_j}$,
- T ranges in the set of d -planes of $\mathbb{R}^p$.

Theorem (Aamari and Levrard 2019)

Picking $h \asymp (\log n/n)^{1/d}$, then with high probability,

$$\max_{1 \leq j \leq n} \angle(T_{X_j} M, \hat{T}_j^{\text{PCA}}) \lesssim \left(\frac{\log n}{n} \right)^{1/d},$$

where $\angle(T, T') := \|\pi_T - \pi_{T'}\|_{\text{op}}$.

Manifold Estimation from Random Sample

Proposition (Aamari and Levrard 2019)

An i.i.d. n -sample $\mathbb{X}_n = \{X_1, \dots, X_n\}$ of $P \in \mathcal{P}_{\text{rch}_{\min}}^{d,D}$ fulfills:

$$\begin{aligned} - \max_{X_j \in \mathbb{X}_n} d(X_j, M) &= 0 & - \sup_{p \in M} d(p, \mathbb{X}_n) &\lesssim (\log n/n)^{1/d}. \end{aligned}$$

The family of d -planes $\hat{T}_{\mathbb{X}_n}^{\text{PCA}}$ built from local PCA fulfills

$$\max_{X_j \in \mathbb{X}_n} \angle(T_{X_j} M, \hat{T}_{X_j}) \lesssim (\log n/n)^{1/d}.$$

$\Rightarrow$ With high probability, we get precision:

$$\varepsilon = d_H(M, \hat{M}_{\text{PATCH}}) \lesssim \left(\frac{\log n}{n} \right)^{2/d}.$$

This rate is minimax optimal

Optimality: Studying the Minimax Risk

The **minimax risk** over the statistical model $\mathcal{P}$ is

$$\inf_{\hat{M}_n} \sup_{P \in \mathcal{P}} \mathbb{E}_{P^n} \left[d_H(M, \hat{M}_n) \right],$$

where $\hat{M}_n = \hat{M}_n(\mathbb{X}_n)$ ranges over all the estimators based on data $\mathbb{X}_n = \{X_1, \dots, X_n\}$.

Optimality: Studying the Minimax Risk

The **minimax risk** over the statistical model $\mathcal{P}$ is

$$\inf_{\hat{M}_n} \sup_{P \in \mathcal{P}} \mathbb{E}_{P^n} \left[d_H(M, \hat{M}_n) \right],$$

where $\hat{M}_n = \hat{M}_n(\mathbb{X}_n)$ ranges over all the estimators based on data $\mathbb{X}_n = \{X_1, \dots, X_n\}$.

Proposition (Genovese et al. 2012b; Kim and Zhou 2015)

For n large enough,

$$\inf_{\hat{M}_n} \sup_{P \in \mathcal{P}} \mathbb{E}_{P^n} \left[d_H(M, \hat{M}_n) \right] \leq C \left(\frac{\log n}{n} \right)^{\frac{2}{d}},$$

where $C = C_{d, \text{rch}_{\min}}$

Optimality: Studying the Minimax Risk

The **minimax risk** over the statistical model $\mathcal{P}$ is

$$\inf_{\hat{M}_n} \sup_{P \in \mathcal{P}} \mathbb{E}_{P^n} \left[d_H(M, \hat{M}_n) \right],$$

where $\hat{M}_n = \hat{M}_n(\mathbb{X}_n)$ ranges over all the estimators based on data $\mathbb{X}_n = \{X_1, \dots, X_n\}$.

Proposition (Genovese et al. 2012b; Kim and Zhou 2015)

For n large enough, (+ mild technical assumptions)

$$c \left(\frac{\log n}{n} \right)^{\frac{2}{d}} \leq \inf_{\hat{M}_n} \sup_{P \in \mathcal{P}} \mathbb{E}_{P^n} \left[d_H(M, \hat{M}_n) \right] \leq C \left(\frac{\log n}{n} \right)^{\frac{2}{d}},$$

where $C = C_{d, \text{rch}_{\min}}$ and $c = c_{\text{rch}_{\min}}$.

Lower Bound Technique: Le Cam's Lemma

Theorem (L. Le Cam)

For all $P_0, P_1 \in \mathcal{P}$,

$$\inf_{\hat{M}_n} \sup_{P \in \mathcal{P}} \mathbb{E}_{P^n} [\mathrm{d}_H(M, \hat{M}_n)] \geq \frac{1}{2} \mathrm{d}_H(M_0, M_1) (1 - \mathrm{TV}(P_0, P_1))^n,$$

where

$$\mathrm{TV}(P_0, P_1) = \sup_{B \in \mathcal{B}(\mathbb{R}^p)} |P_0(B) - P_1(B)|$$

denotes the total variation distance between P_0 and P_1 .

Lower Bound Technique: Le Cam's Lemma

Theorem (L. Le Cam)

For all $P_0, P_1 \in \mathcal{P}$,

$$\inf_{\hat{M}_n} \sup_{P \in \mathcal{P}} \mathbb{E}_{P^n} [\mathrm{d}_H(M, \hat{M}_n)] \geq \frac{1}{2} \mathrm{d}_H(M_0, M_1) (1 - \mathrm{TV}(P_0, P_1))^n,$$

where

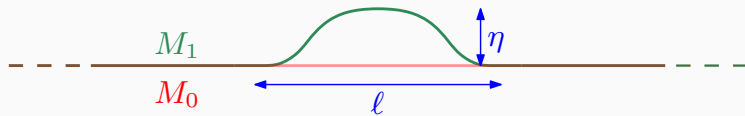
$$\mathrm{TV}(P_0, P_1) = \sup_{B \in \mathcal{B}(\mathbb{R}^p)} |P_0(B) - P_1(B)|$$

denotes the total variation distance between P_0 and P_1 .

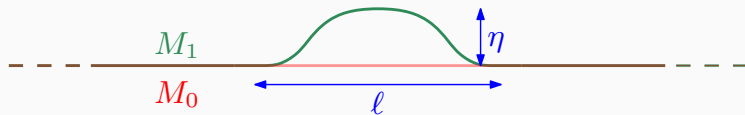
Deriving a good lower bound amounts to find P_0, P_1 such that:

- $P_0, P_1 \in \mathcal{P}$,
- $\mathrm{d}_H(M_0, M_1)$ is large,
- $\mathrm{TV}(P_0, P_1)$ is small.

Le Cam's Lemma Heuristic

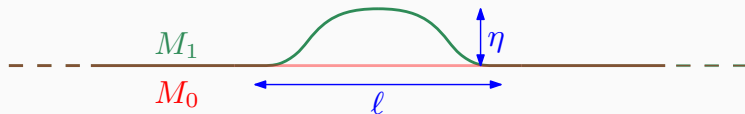


Le Cam's Lemma Heuristic



- P_0 and P_1 both belong to $\mathcal{P}$ as soon as $\eta \lesssim \ell^2$,
- $d_H(M_0, M_1) \geq \eta$,
- $TV(P_0, P_1) \lesssim \ell^d$.

Le Cam's Lemma Heuristic



- P_0 and P_1 both belong to $\mathcal{P}$ as soon as $\eta \lesssim \ell^2$,
- $d_H(M_0, M_1) \geq \eta$,
- $\text{TV}(P_0, P_1) \lesssim \ell^d$.

Hence, for $\eta \approx \ell^2$ and $\ell \approx (1/n)^{1/d}$,

$$\inf_{\hat{M}_n} \sup_{P \in \mathcal{P}} \mathbb{E}_{P^n} [d_H(M, \hat{M}_n)] \gtrsim \eta (1 - \ell^d)^n \approx \ell^2 (1 - \ell^d)^n \approx (1/n)^{2/d}.$$

What if the Curve isn't Closed?

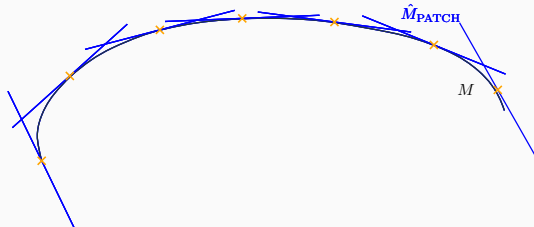
Perform local PCA at each point $X_j \in \mathbb{X}_n$:

$$\hat{T}_j \in \arg \min_{T \in \mathbb{G}^{D,d}} P_n^{(j)} \left[\|x - \pi_T(x)\|^2 \mathbf{1}_{B(0,h)}(x) \right],$$

and take

$$\hat{M}_{\text{PATCH}} := \bigcup_{j=1}^n B_{\hat{T}_j}(0, h).$$

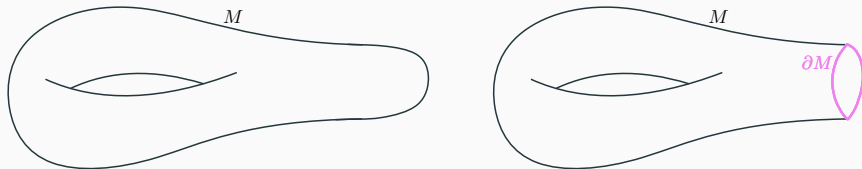
- + Local PCA still estimates tangent spaces up to angle $\lesssim (\log n/n)^{1/d}$.
- Nearby "boundary points", the patches extend too far away from M .



Boundary Manifold Model

We let $\mathcal{P}^\partial := \mathcal{P}_{\text{rch}_{\min}, \text{rch}_{\partial, \min}}^{d, D}$ denote the set of distributions P over $\mathbb{R}^p$ such that

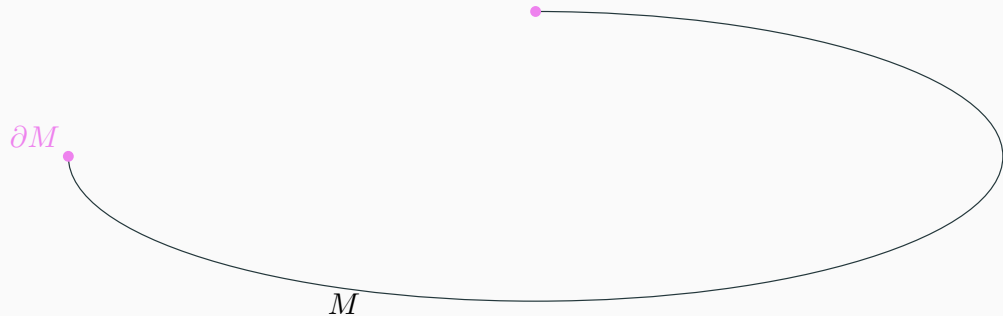
- Its support $M = \text{supp}(P) \subset \mathbb{R}^p$ satisfies:
 - M is a $\mathcal{C}^2$ submanifold **with boundary**;
 - M has reach bounded away from zero $\text{rch}_M \geq \text{rch}_{\min} > 0$;
 - ∂M has reach bounded away from zero $\text{rch}_{\partial M} \geq \text{rch}_{\partial, \min} > 0$.



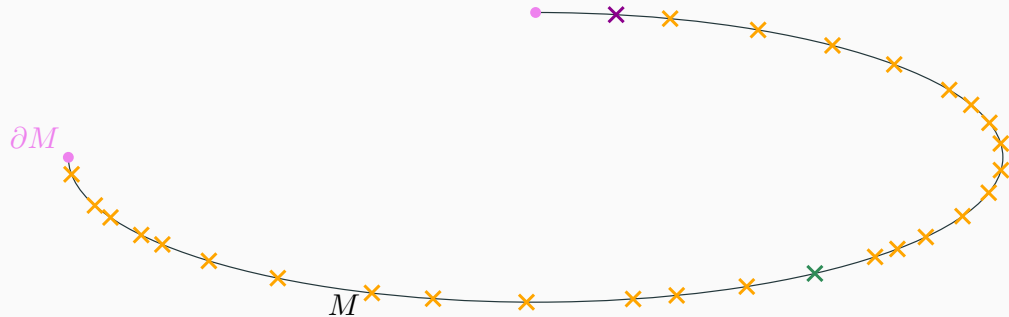
- P is roughly uniform on M :

$$f = dP/d\text{vol}_M \text{ exists and } f_{\min} \leq f \leq f_{\max}.$$

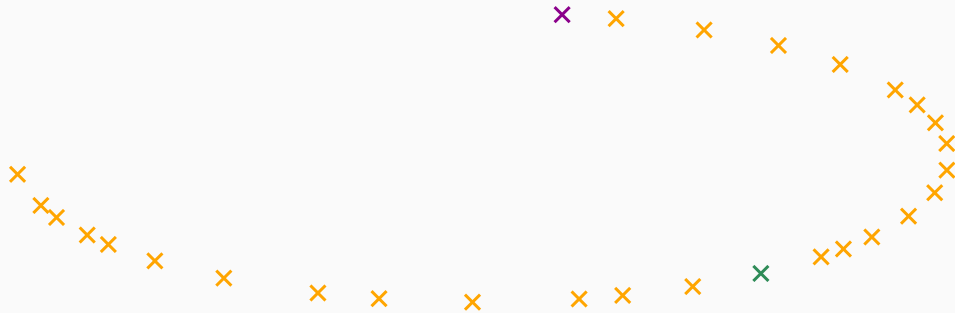
Insights on Boundary Point Detection



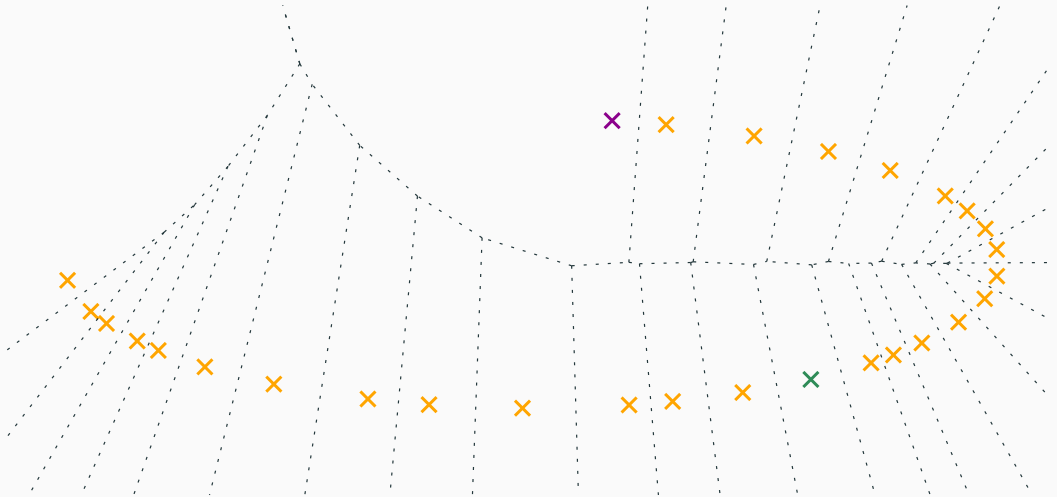
Insights on Boundary Point Detection



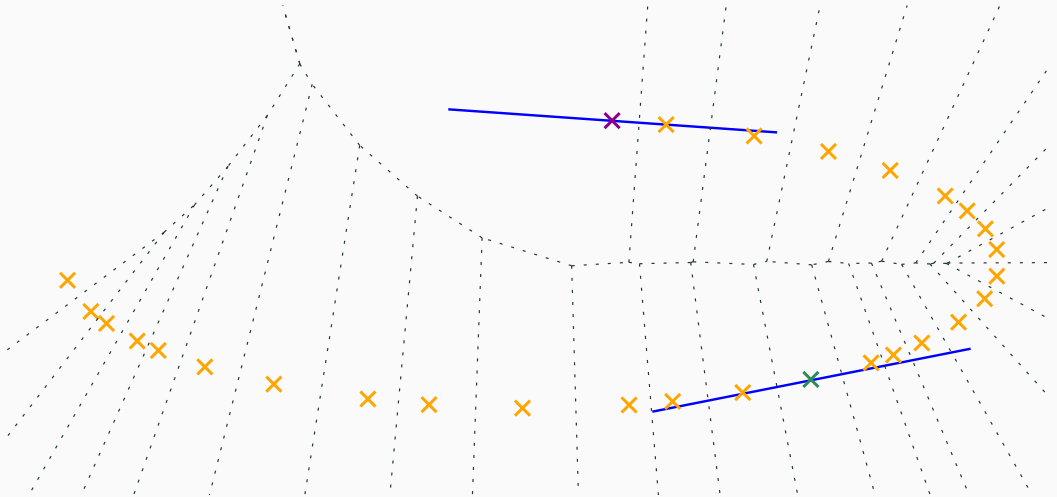
Insights on Boundary Point Detection



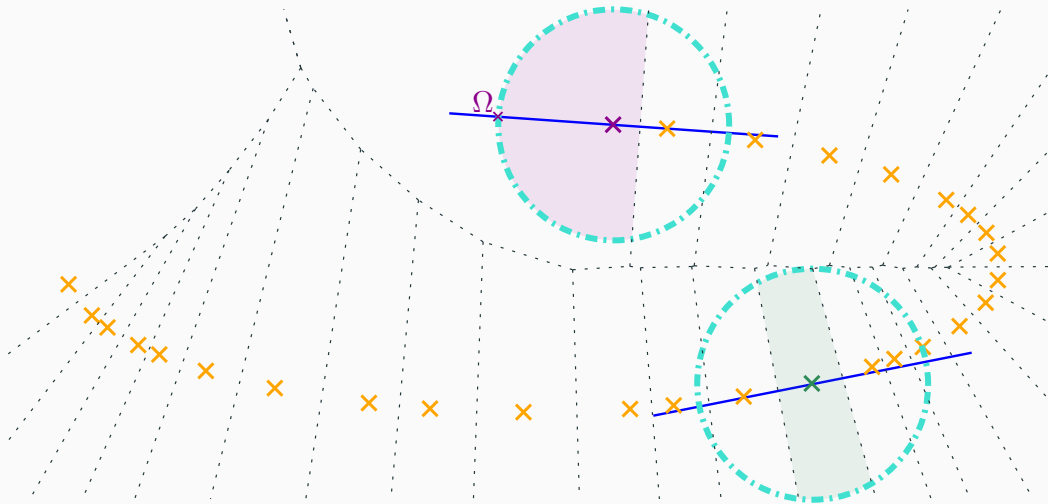
Insights on Boundary Point Detection



Insights on Boundary Point Detection



Insights on Boundary Point Detection



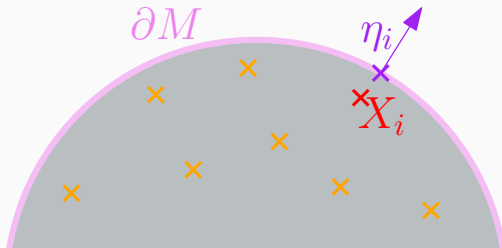
Boundary Observations: Definition

Write

$$\text{Vor}_{R_0}^{(j)}(X_i) := \left\{ O \in \hat{T}_j \mid \mathring{B}(O, \|O - \pi_{\hat{T}_j}(X_i - X_j)\|) \cap \pi_{\hat{T}_j}(B(X_j, R_0) \cap \mathbb{X}_n - X_j) = \emptyset \right\}.$$

Define the set of *boundary observations* as

$$\mathcal{Y}_{R_0, r, \rho} := \left\{ X_i \in \mathbb{X}_n \mid \exists X_j \in B(X_i, r) \cap \mathbb{X}_n \text{ s.t. } \text{Diam}(\text{Vor}_{R_0}^{(j)}(X_i)) \geq \rho \right\}.$$



Boundary Observations: Definition

Write

$$\text{Vor}_{R_0}^{(j)}(X_i) := \left\{ O \in \hat{T}_j \mid \mathring{B}(O, \|O - \pi_{\hat{T}_j}(X_i - X_j)\|) \cap \pi_{\hat{T}_j}(B(X_j, R_0) \cap \mathbb{X}_n - X_j) = \emptyset \right\}.$$

Define the set of *boundary observations* as

$$\mathcal{Y}_{R_0, r, \rho} := \left\{ X_i \in \mathbb{X}_n \mid \exists X_j \in B(X_i, r) \cap \mathbb{X}_n \text{ s.t. } \text{Diam}(\text{Vor}_{R_0}^{(j)}(X_i)) \geq \rho \right\}.$$



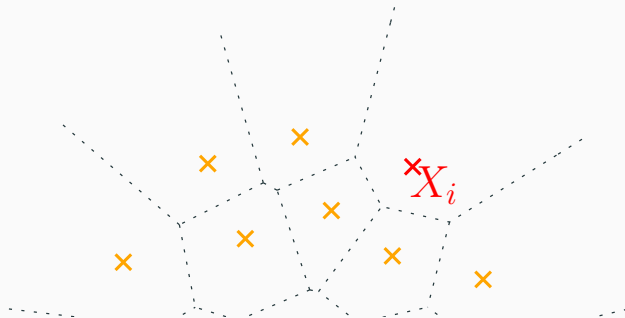
Boundary Observations: Definition

Write

$$\text{Vor}_{R_0}^{(j)}(X_i) := \left\{ O \in \hat{T}_j \mid \mathring{B}(O, \|O - \pi_{\hat{T}_j}(X_i - X_j)\|) \cap \pi_{\hat{T}_j}(\text{B}(X_j, R_0) \cap \mathbb{X}_n - X_j) = \emptyset \right\}.$$

Define the set of *boundary observations* as

$$\mathcal{Y}_{R_0, r, \rho} := \left\{ X_i \in \mathbb{X}_n \mid \exists X_j \in \text{B}(X_i, r) \cap \mathbb{X}_n \text{ s.t. } \text{Diam}(\text{Vor}_{R_0}^{(j)}(X_i)) \geq \rho \right\}.$$



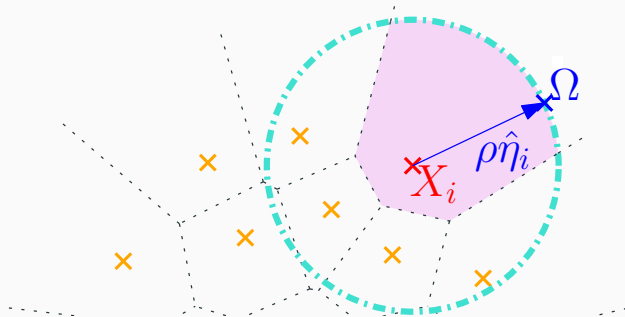
Boundary Observations: Definition

Write

$$\text{Vor}_{R_0}^{(j)}(X_i) := \left\{ O \in \hat{T}_j \mid \mathring{B}(O, \|O - \pi_{\hat{T}_j}(X_i - X_j)\|) \cap \pi_{\hat{T}_j}(B(X_j, R_0) \cap \mathbb{X}_n - X_j) = \emptyset \right\}.$$

Define the set of *boundary observations* as

$$\mathcal{Y}_{R_0, r, \rho} := \left\{ X_i \in \mathbb{X}_n \mid \exists X_j \in B(X_i, r) \cap \mathbb{X}_n \text{ s.t. } \text{Diam}(\text{Vor}_{R_0}^{(j)}(X_i)) \geq \rho \right\}.$$



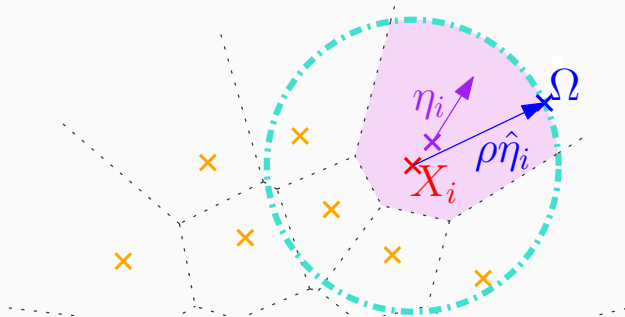
Boundary Observations: Definition

Write

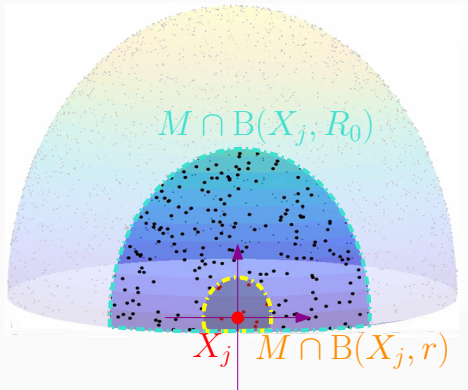
$$\text{Vor}_{R_0}^{(j)}(X_i) := \left\{ O \in \hat{T}_j \mid \mathring{B}(O, \|O - \pi_{\hat{T}_j}(X_i - X_j)\|) \cap \pi_{\hat{T}_j}(\text{B}(X_j, R_0) \cap \mathbb{X}_n - X_j) = \emptyset \right\}.$$

Define the set of *boundary observations* as

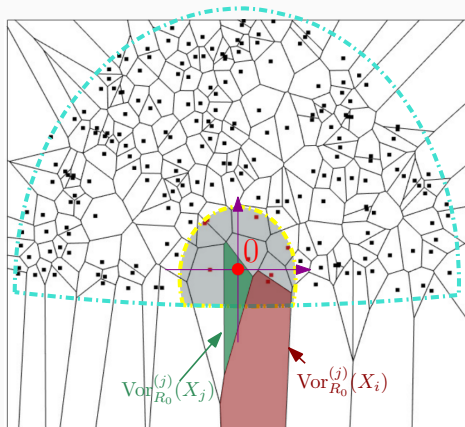
$$\mathcal{Y}_{R_0, r, \rho} := \left\{ X_i \in \mathbb{X}_n \mid \exists X_j \in \text{B}(X_i, r) \cap \mathbb{X}_n \text{ s.t. } \text{Diam}(\text{Vor}_{R_0}^{(j)}(X_i)) \geq \rho \right\}.$$



Boundary Observations: Illustration



$\pi_{\hat{T}_j}$



Guarantees for Boundary Detection and Normals

Choosing the parameters properly, we have the following with high probability:

If $\partial M = \emptyset$, then $\mathcal{Y}_{R_0, r, \rho} = \emptyset$;

If $\partial M \neq \emptyset$ then:

For all $X_i \in \mathcal{Y}_{R_0, r, \rho}$,

$$d(X_i, \partial M) \lesssim \left(\frac{\log n}{n} \right)^{\frac{2}{d+1}}.$$

For all $x \in \partial M$,

$$d(x, \mathcal{Y}_{R_0, r, \rho}) \lesssim \left(\frac{\log n}{n} \right)^{\frac{1}{d+1}}.$$

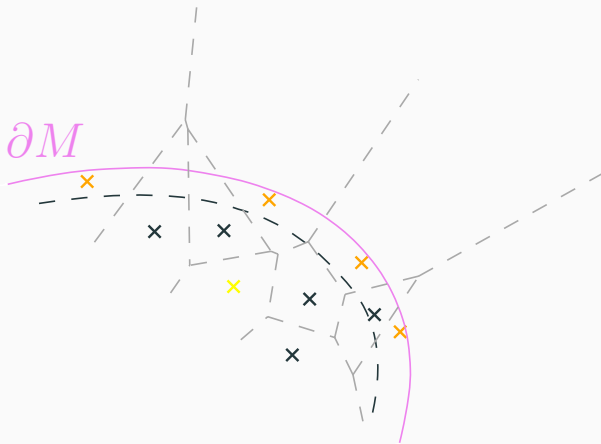
For all $X_i \in \mathcal{Y}_{R_0, r, \rho}$,

$$\|\eta_{\pi_{\partial M}(X_i)} - \tilde{\eta}_i\| \lesssim \left(\frac{\log n}{n} \right)^{\frac{1}{d+1}}.$$

Guarantees: Illustration

Write

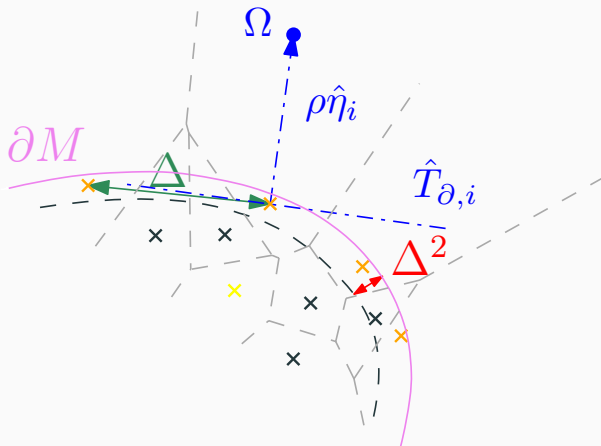
$$\Delta := \left(\frac{\log n}{n} \right)^{\frac{1}{d+1}}.$$



Guarantees: Illustration

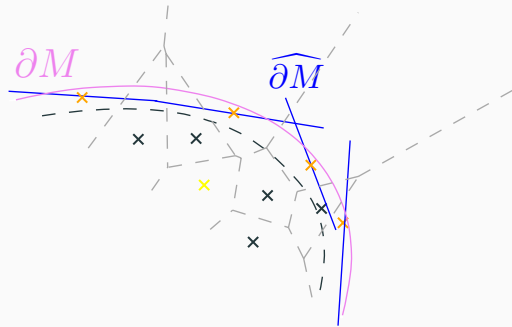
Write

$$\Delta := \left(\frac{\log n}{n} \right)^{\frac{1}{d+1}}.$$



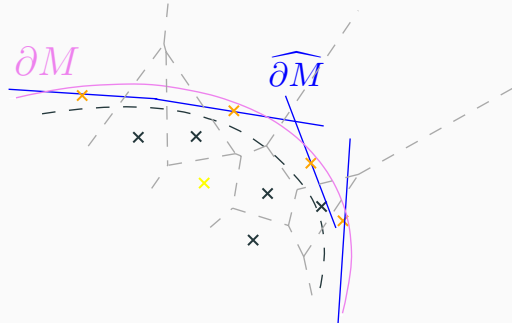
Boundary Estimation

Boundary points $\mathcal{Y}$
Boundary's tangents estimates $\hat{T}_{\partial,i}$ } $\Rightarrow$ local linear patches $\widehat{\partial M}$



Boundary Estimation

Boundary points $\mathcal{Y}$
Boundary's tangents estimates $\hat{T}_{\partial,i}$ } $\Rightarrow$ local linear patches $\widehat{\partial M}$



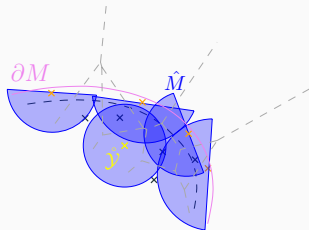
Theorem (Aamari, Aaron, Levrard – 2023)

$$\mathbb{E}[\mathrm{d}_{\mathrm{H}}(\partial M, \widehat{\partial M})] \lesssim \left(\frac{\log n}{n} \right)^{\frac{2}{d+1}}$$

(minimax optimal over $\mathcal{P}_{\mathrm{rch}_{\min}, \mathrm{rch}_{\partial, \min}}^{d,D}$)

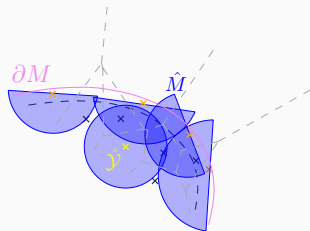
Estimation with Boundary

Boundary & Interior points $\mathcal{Y}$ & $\mathbb{X}_n \setminus \mathcal{Y}$
Boundary's tangents estimates $\hat{T}_{\partial,i}$
Manifold's tangents estimates $\hat{T}_i$ } $\Rightarrow$ local linear patches / half-patches $\hat{M}$



Estimation with Boundary

Boundary & Interior points $\mathcal{Y}$ & $\mathbb{X}_n \setminus \mathcal{Y}$
Boundary's tangents estimates $\hat{T}_{\partial,i}$
Manifold's tangents estimates $\hat{T}_i$ } $\Rightarrow$ local linear patches / half-patches $\hat{M}$

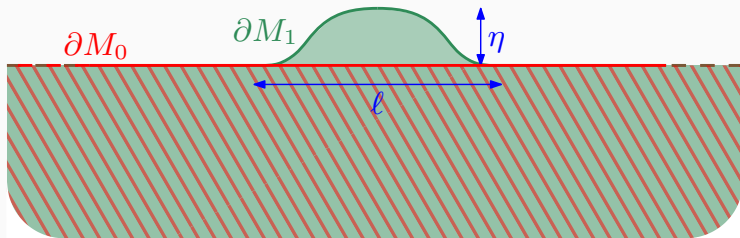


Theorem (Aamari, Aaron, and Levrard 2023)

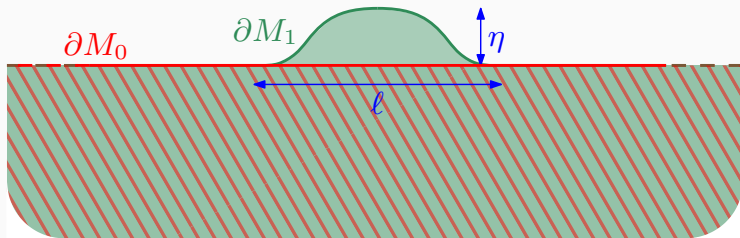
$$\mathbb{E}[\mathrm{d}_H(M, \widehat{M})] \lesssim \left(\frac{\log n}{n} \right)^{\frac{2}{d+1}}$$

(minimax optimal over $\mathcal{P}_{\mathrm{rch}_{\min}, \mathrm{rch}_{\partial, \min}}^{d,D}$)

Le Cam's Lemma Heuristic

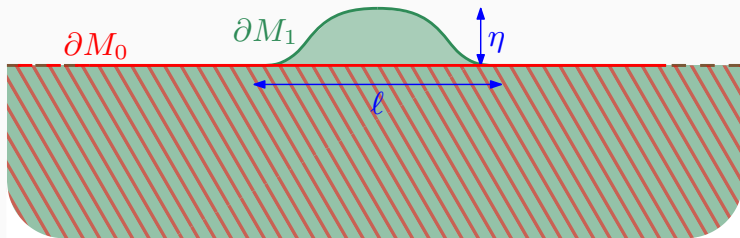


Le Cam's Lemma Heuristic



- P_0 and P_1 both belong to $\mathcal{P}^\partial$ as soon as $\eta \lesssim \ell^2$,
- $d_H(M_0, M_1) \geq \eta$,
- $\text{TV}(P_0, P_1) \lesssim \ell^{d-1} \eta$.

Le Cam's Lemma Heuristic

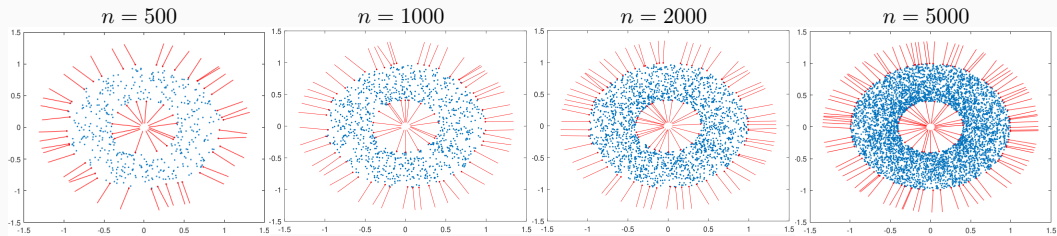


- P_0 and P_1 both belong to $\mathcal{P}^\partial$ as soon as $\eta \lesssim \ell^2$,
- $d_H(M_0, M_1) \geq \eta$,
- $TV(P_0, P_1) \lesssim \ell^{d-1}\eta$.

Hence, for $\eta \approx \ell^2$ and $\ell \approx (1/n)^{1/(d+1)}$,

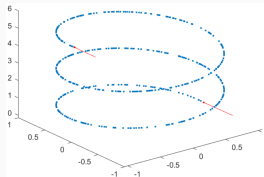
$$\inf_{\hat{M}_n} \sup_{P \in \mathcal{P}} \mathbb{E}_{P^n} [d_H(M, \hat{M}_n)] \gtrsim \eta (1 - \ell^{d-1}\eta)^n \approx \ell^2 (1 - \ell^{d+1})^n \approx (1/n)^{2/(d+1)}.$$

Annulus

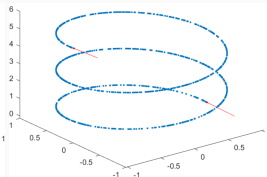


Spiral

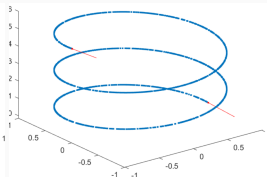
$n = 500$



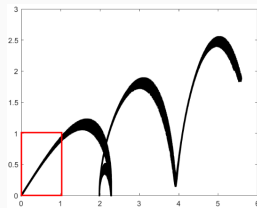
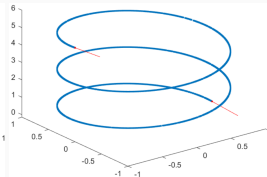
$n = 1000$



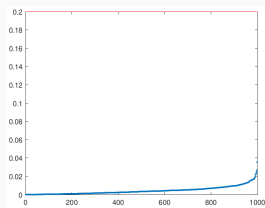
$n = 2000$



$n = 5000$



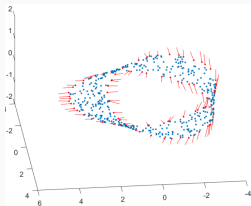
Calibration of R_0



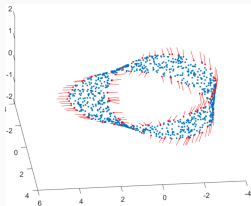
Calibration of ρ

Möbius strip

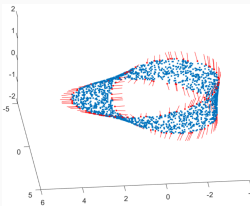
$n = 500$



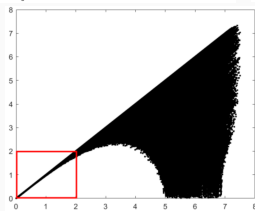
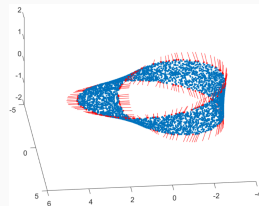
$n = 1000$



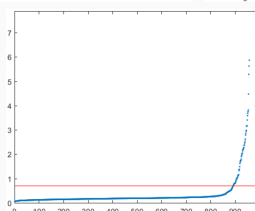
$n = 2000$



$n = 5000$



Calibration of R_0



Calibration of ρ

Influence of noise

Additive Noise

Crucial limitation: If significant noise is added, all the above methods fail!

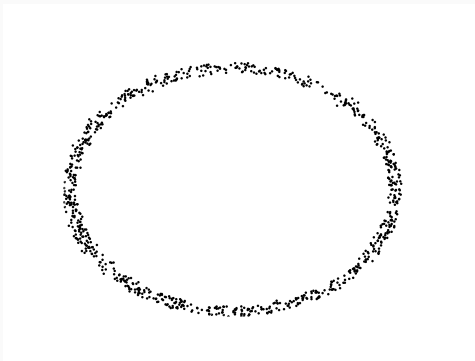


Figure 1: Circle with additive noise amplitude $\sigma > 0$

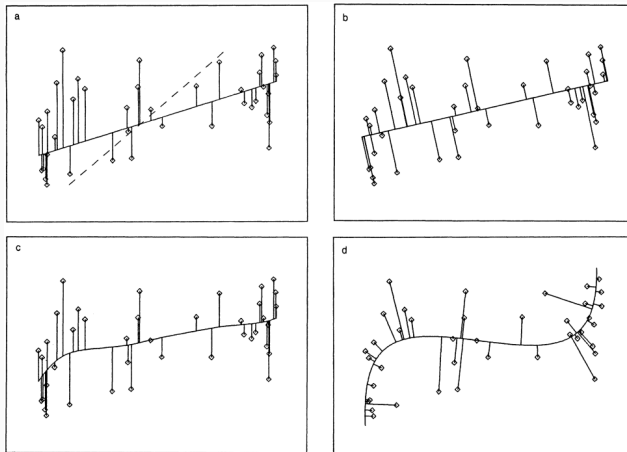


Figure 1. (a) The linear regression line minimizes the sum of squared deviations in the response variable. (b) The principal-component line minimizes the sum of squared deviations in all of the variables. (c) The smooth regression curve minimizes the sum of squared deviations in the response variable, subject to smoothness constraints. (d) The principal curve minimizes the sum of squared deviations in all of the variables, subject to smoothness constraints.

A Zoology of Noise Models

As opposed to nonparametric regression, many natural noise models:

- $Y = X + \varepsilon$ with $X \in M$
and $X \perp\!\!\!\perp \varepsilon \in \mathbb{R}^p$ such that $\mathbb{E}[X] = 0$ (Convolution)
Fefferman et al. 2023; Genovese et al. 2012b
- $Y = X + \varepsilon$ with $X \in M$
and $\varepsilon \in (T_X M)^\perp$ such that $\mathbb{E}[\varepsilon|X] = 0$ (Orthogonal noise)
Genovese et al. 2012a; Puchkin and Spokoiny 2022
- $Y \sim \text{Unif}_{M^\sigma}$ (Ambient uniform)
Aizenbud and Sober 2021

A Zoology of Noise Models

As opposed to nonparametric regression, many natural noise models:

- $Y = X + \varepsilon$ with $X \in M$
and $X \perp\!\!\!\perp \varepsilon \in \mathbb{R}^p$ such that $\mathbb{E}[X] = 0$ (Convolution)
Fefferman et al. 2023; Genovese et al. 2012b
- $Y = X + \varepsilon$ with $X \in M$
and $\varepsilon \in (T_X M)^\perp$ such that $\mathbb{E}[\varepsilon|X] = 0$ (Orthogonal noise)
Genovese et al. 2012a; Puchkin and Spokoiny 2022
- $Y \sim \text{Unif}_{M^\sigma}$ (Ambient uniform)
Aizenbud and Sober 2021

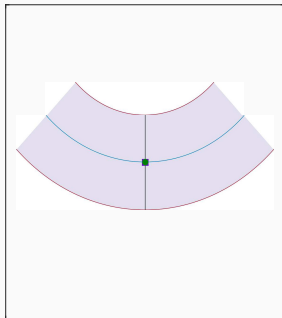
Take away:

Minimax rates for manifold estimation in the presence of noise are not fully understood.

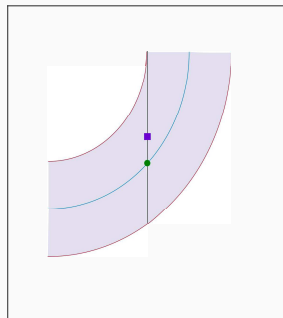
About Noise not Being Centered

• $p \in \mathcal{M}$

■ $\mathbb{E}[Y|X = x]$



(a)



(b)

Figure 3: From Aizenbud and Sober 2021

Problem

An error in the tangent space yield apparent noise not centered, whatever the type of noise.

Alternating local PCA and Nonparametric Regression

Iterative algorithm

Aizenbud and Sober 2021; Puchkin and Spokoiny 2022

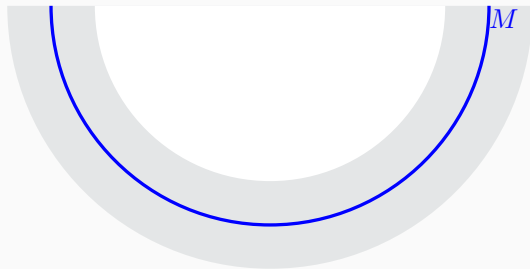
Tangents Initialize local coordinates with local PCA at scale $h_0 \simeq 1$.

Denoising In these coordinates, apply classical nonparametric regression at scale $h_1 < h_0$.

Tangents Store these new local coordinates and associated denoised points

Denoising In these coordinates, apply classical nonparametric regression at scale $h_2 < h_1$.

...



Alternating local PCA and Nonparametric Regression

Iterative algorithm

Aizenbud and Sober 2021; Puchkin and Spokoiny 2022

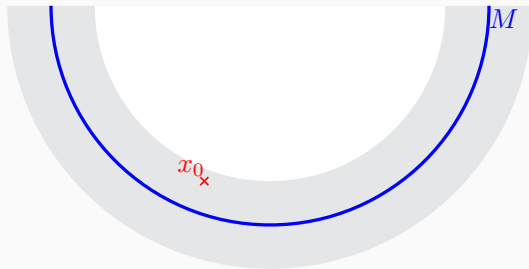
Tangents Initialize local coordinates with local PCA at scale $h_0 \simeq 1$.

Denoising In these coordinates, apply classical nonparametric regression at scale $h_1 < h_0$.

Tangents Store these new local coordinates and associated denoised points

Denoising In these coordinates, apply classical nonparametric regression at scale $h_2 < h_1$.

...



Alternating local PCA and Nonparametric Regression

Iterative algorithm

Aizenbud and Sober 2021; Puchkin and Spokoiny 2022

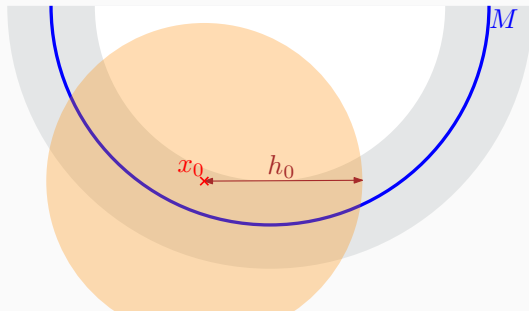
Tangents Initialize local coordinates with local PCA at scale $h_0 \simeq 1$.

Denoising In these coordinates, apply classical nonparametric regression at scale $h_1 < h_0$.

Tangents Store these new local coordinates and associated denoised points

Denoising In these coordinates, apply classical nonparametric regression at scale $h_2 < h_1$.

...



Alternating local PCA and Nonparametric Regression

Iterative algorithm

Aizenbud and Sober 2021; Puchkin and Spokoiny 2022

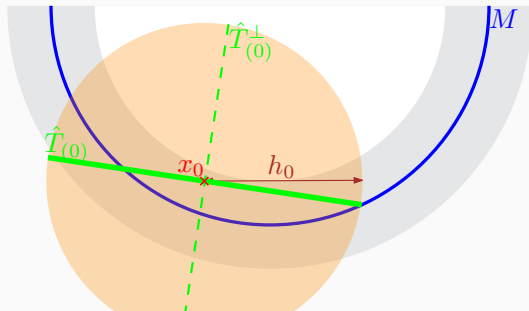
Tangents Initialize local coordinates with local PCA at scale $h_0 \simeq 1$.

Denoising In these coordinates, apply classical nonparametric regression at scale $h_1 < h_0$.

Tangents Store these new local coordinates and associated denoised points

Denoising In these coordinates, apply classical nonparametric regression at scale $h_2 < h_1$.

...



Alternating local PCA and Nonparametric Regression

Iterative algorithm

Aizenbud and Sober 2021; Puchkin and Spokoiny 2022

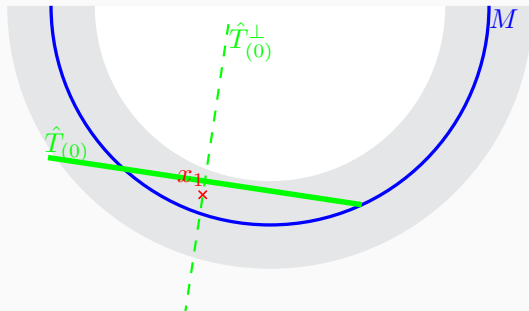
Tangents Initialize local coordinates with local PCA at scale $h_0 \simeq 1$.

Denoising In these coordinates, apply classical nonparametric regression at scale $h_1 < h_0$.

Tangents Store these new local coordinates and associated denoised points

Denoising In these coordinates, apply classical nonparametric regression at scale $h_2 < h_1$.

...



Alternating local PCA and Nonparametric Regression

Iterative algorithm

Aizenbud and Sober 2021; Puchkin and Spokoiny 2022

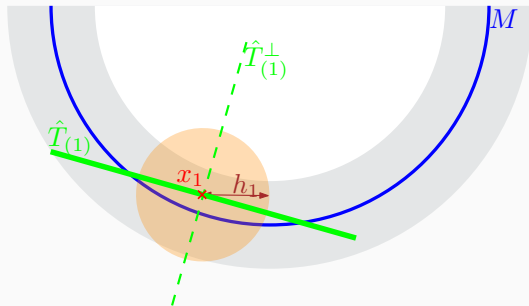
Tangents Initialize local coordinates with local PCA at scale $h_0 \simeq 1$.

Denoising In these coordinates, apply classical nonparametric regression at scale $h_1 < h_0$.

Tangents Store these new local coordinates and associated denoised points

Denoising In these coordinates, apply classical nonparametric regression at scale $h_2 < h_1$.

...



Alternating local PCA and Nonparametric Regression

Iterative algorithm

Aizenbud and Sober 2021; Puchkin and Spokoiny 2022

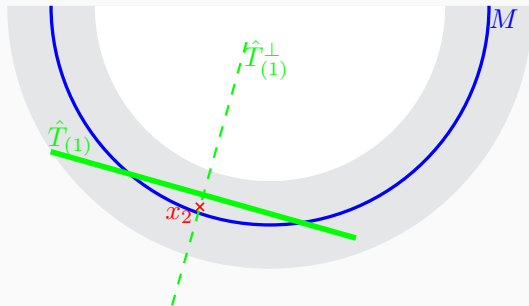
Tangents Initialize local coordinates with local PCA at scale $h_0 \simeq 1$.

Denoising In these coordinates, apply classical nonparametric regression at scale $h_1 < h_0$.

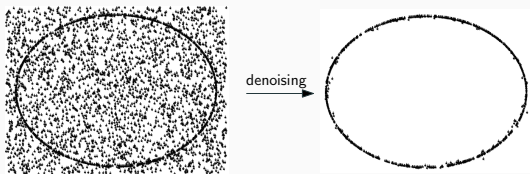
Tangents Store these new local coordinates and associated denoised points

Denoising In these coordinates, apply classical nonparametric regression at scale $h_2 < h_1$.

...



$$\mathcal{P}_{\beta, Q_0}^{(\text{clutter})} = \{\beta P + (1 - \beta)Q_0, P \in \mathcal{P}\}.$$



Theorem (Aamari, Levrard – 2018)

With a decluttering procedure removing points from $Q_0 = \text{Unif}_{B(0,R)}$, we can build an estimator such that

$$\sup_{P \in \mathcal{P}^{(\text{clutter})}} \mathbb{E}_{P^n} \left[d_H(M, \hat{M}_n) \right] \lesssim \left(\frac{\log n}{n} \right)^{\frac{2}{d}}.$$

Remark: This procedure may fail for other Q_0 's, even if Q_0 is known.

Parameter Selection

r -Convex Hull

For all $t \geq 0$, the t -convex hull of $A \subset \mathbb{R}^p$ is

$$\text{Conv}_t(A) := \bigcup_{\substack{\sigma \subset A \\ \text{rad}(\sigma) \leq t}} \text{Conv}(\sigma),$$

where $\text{rad}(\sigma)$ is the radius of the smallest ball enclosing σ .



Figure 4: from Vicent Divol's PhD Defense

r -Convex Hull

For all $t \geq 0$, the t -convex hull of $A \subset \mathbb{R}^p$ is

$$\text{Conv}_t(A) := \bigcup_{\substack{\sigma \subset A \\ \text{rad}(\sigma) \leq t}} \text{Conv}(\sigma),$$

where $\text{rad}(\sigma)$ is the radius of the smallest ball enclosing σ .

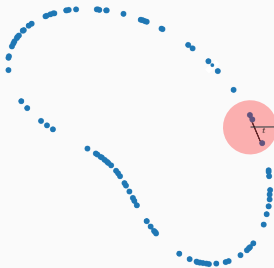


Figure 4: from Vicent Divol's PhD Defense

r -Convex Hull

For all $t \geq 0$, the t -convex hull of $A \subset \mathbb{R}^p$ is

$$\text{Conv}_t(A) := \bigcup_{\substack{\sigma \subset A \\ \text{rad}(\sigma) \leq t}} \text{Conv}(\sigma),$$

where $\text{rad}(\sigma)$ is the radius of the smallest ball enclosing σ .

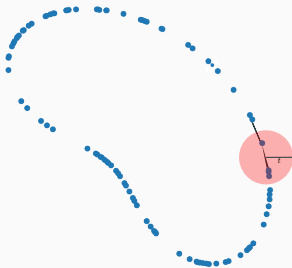


Figure 4: from Vicent Divol's PhD Defense

r -Convex Hull

For all $t \geq 0$, the t -convex hull of $A \subset \mathbb{R}^p$ is

$$\text{Conv}_t(A) := \bigcup_{\substack{\sigma \subset A \\ \text{rad}(\sigma) \leq t}} \text{Conv}(\sigma),$$

where $\text{rad}(\sigma)$ is the radius of the smallest ball enclosing σ .

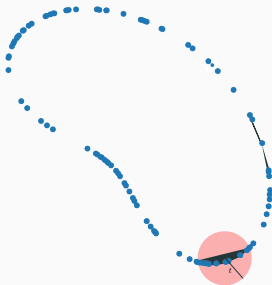


Figure 4: from Vicent Divol's PhD Defense

r -Convex Hull

For all $t \geq 0$, the t -convex hull of $A \subset \mathbb{R}^p$ is

$$\text{Conv}_t(A) := \bigcup_{\substack{\sigma \subset A \\ \text{rad}(\sigma) \leq t}} \text{Conv}(\sigma),$$

where $\text{rad}(\sigma)$ is the radius of the smallest ball enclosing σ .

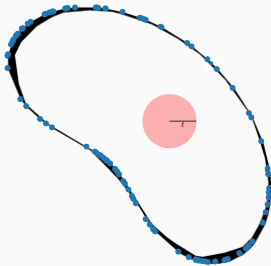


Figure 4: from Vicent Divol's PhD Defense

Reconstruction from r -Convex Hull

$$\text{Let } t^*(A) := \inf \{t < \text{rch}_M \mid \pi_M(\text{Conv}_t(A)) = M\}$$

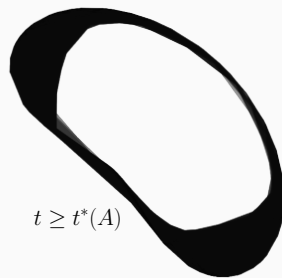
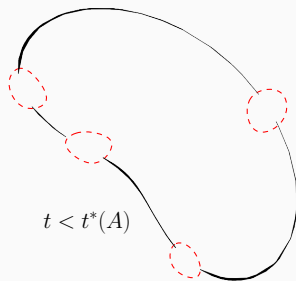


Figure 5: from Vicent Divol's PhD Defense

Reconstruction from r -Convex Hull

$$\text{Let } t^*(A) := \inf \{t < \text{rch}_M \mid \pi_M(\text{Conv}_t(A)) = M\}$$

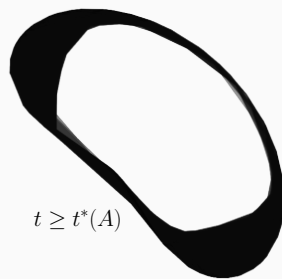
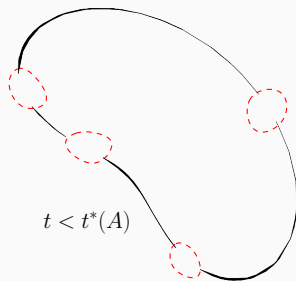


Figure 5: from Vicent Divol's PhD Defense

$\Rightarrow$ To reconstruct, need to pick $t > t^*(A)$ but as small as possible.

Reconstruction from r -Convex Hull

Theorem (Divol 2021)

There exists $C = C_{\mathcal{P}} > 0$ such that picking $t = C (\log n/n)^{1/d}$, then for all $P \in \mathcal{P}$ and $n \geq 1$ large enough,

$$d_H(M, \text{Conv}_t(\mathcal{X}_n)) \lesssim \left(\frac{\log n}{n} \right)^{2/d}.$$

Reconstruction from r -Convex Hull

Theorem (Divol 2021)

There exists $C = C_{\mathcal{P}} > 0$ such that picking $t = C (\log n/n)^{1/d}$, then for all $P \in \mathcal{P}$ and $n \geq 1$ large enough,

$$d_H(M, \text{Conv}_t(\mathcal{X}_n)) \lesssim \left(\frac{\log n}{n} \right)^{2/d}.$$

Limitation

In practice, need to **calibrate** the constant C .

(or equivalently t)

Idea

Compare each estimator $\text{Conv}_t(\mathcal{X}_n)$ with the **most overfitting** one $\text{Conv}_t(\mathcal{X}_n) = \mathcal{X}_n$ of the family.

Convexity Defect Function

The **convexity defect function** of $A \subset \mathbb{R}^p$ at scale $t \geq 0$ is

$$h(t, A) := d_H(A, \text{Conv}_t(A))$$

Convexity Defect Function

The **convexity defect function** of $A \subset \mathbb{R}^p$ at scale $t \geq 0$ is

$$h(t, A) := d_H(A, \text{Conv}_t(A))$$

If $\text{rch}_M > 0$, then $h(t, M) \leq t^2/\text{rch}_M$

For point clouds $A = \mathcal{X}_n$, the behavior looks like this:

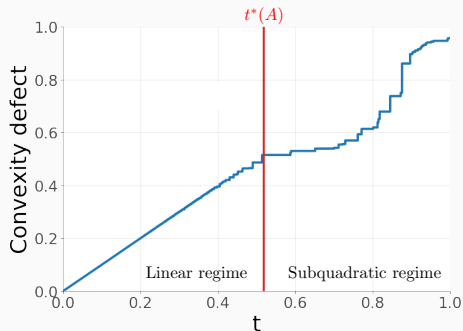
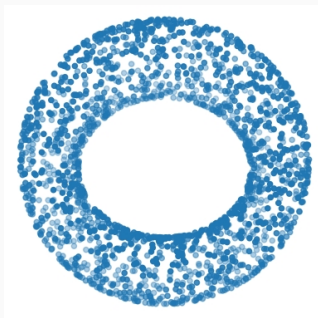


Figure 6: from Vicent Divil's PhD Defense

Scale Parameter Choice

Given $0 < \lambda \leq 1$, define

$$t_\lambda(A) := \inf\{t \in \text{Rad}(A) \mid h(t, A) \leq \lambda t\},$$

where $\text{Rad}(A) = \{\text{rad}(\sigma)\}_{\sigma \subset A}$.

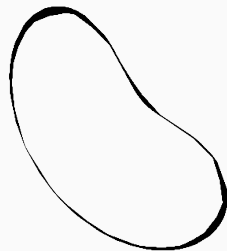
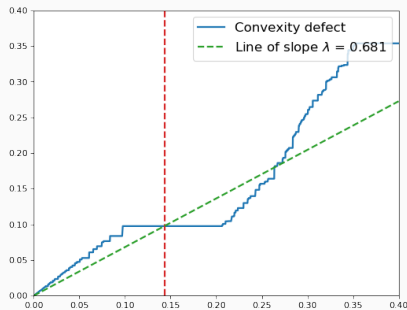


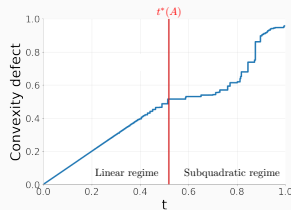
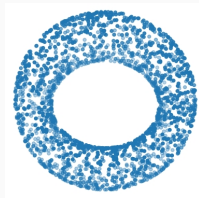
Figure 7: from Vicent Divol's PhD Defense

Scale-Free Manifold Estimation

Theorem (Divol 2021)

Uniformly over $\mathcal{P}$, for all $n \geq 1$ large enough,

$$\mathbb{E}_{P^n} \left[d_H(M, \text{Conv}_{t_\lambda(\mathcal{X}_n)}(\mathcal{X}_n)) \right] \lesssim \left(\frac{\log n}{n} \right)^{2/d}.$$

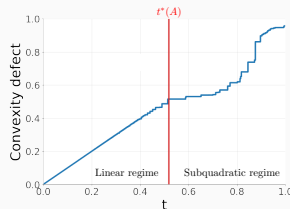
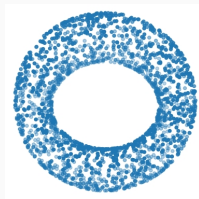


Scale-Free Manifold Estimation

Theorem (Divol 2021)

Uniformly over $\mathcal{P}$, for all $n \geq 1$ large enough,

$$\mathbb{E}_{P^n} [\mathrm{d}_H(M, \mathrm{Conv}_{t_\lambda(\mathcal{X}_n)}(\mathcal{X}_n))] \lesssim \left(\frac{\log n}{n} \right)^{2/d}.$$



Remark: This method is *not* fully parameter-free: choice of $\lambda \geq 1$.

Yet, $\lambda = 1/\sqrt{2}$ works (theoretically) for any dimension $d \geq 1$.

Smoother Manifolds

More Regularity

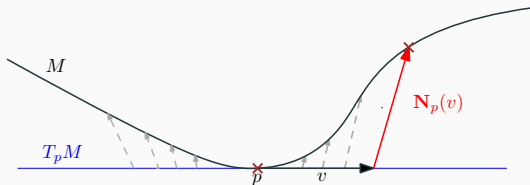
Definition ($\mathcal{C}^2$ Regularity Class)

Submanifolds $M \in \mathcal{C}_{\text{rch}_{\min}}^2$ have local parametrizations

$$\Psi_p : T_p M \longrightarrow M \subset \mathbb{R}^p$$

$$v \longmapsto p + v + \mathbf{N}_p(v)$$

where $\mathbf{N}_p(0) = 0$, $d_0 \mathbf{N}_p = 0$ and $\|d_v \mathbf{N}_p\|_{op} \leq \|v\|/(2\text{rch}_{\min})$.



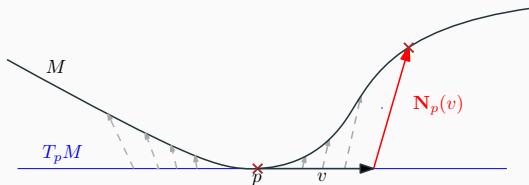
More Regularity

Definition ($\mathcal{C}^k$ Regularity Class, $k \geq 3$)

Let $\mathbf{L} = (L_2, L_3, \dots, L_k)$, and define $\mathcal{C}_{\text{rch}_{\min}, \mathbf{L}}^k$ to be the subset of elements $M \in \mathcal{C}_{\text{rch}_{\min}}^2$ that have local parametrizations

$$\begin{aligned}\Psi_p : T_p M &\longrightarrow M \subset \mathbb{R}^p \\ v &\longmapsto p + v + \mathbf{N}_p(v)\end{aligned}$$

where $\mathbf{N}_p(0) = 0$, $d_0 \mathbf{N}_p = 0$ and $\|d_v^i \mathbf{N}_p\|_{op} \leq L_i$ for $2 \leq i \leq k$.



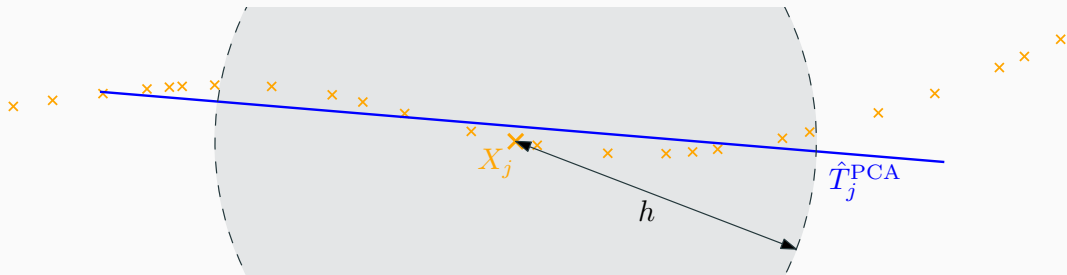
Local PCA

Recall that $P_n^{(j)} = \frac{1}{n} \sum_{\ell \neq j} \delta_{X_\ell - X_j}$, and

$$\hat{T}_j^{\text{PCA}} \in \arg \min_{T \in \mathbb{G}^{D,d}} P_n^{(j)} [\|x - \pi_T(x)\|^2 \mathbb{I}\{B(0, h)\}(x)] .$$

$\mathbb{G}^{p,d}$: space of d -dimensional linear subspaces of $\mathbb{R}^p$;

π_T : orthogonal projection onto T .



Local Polynomials

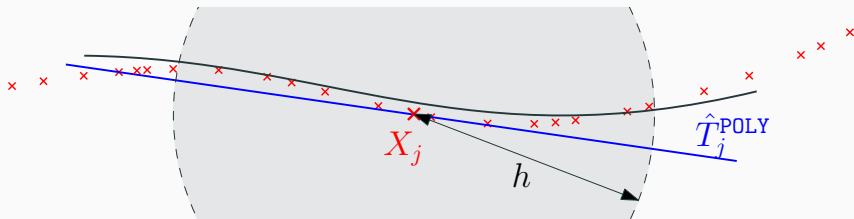
Define $(\hat{T}_j^{\text{POLY}}, \hat{T}_{2,j}, \dots, \hat{T}_{k-1,j})$ to be a minimizer of

$$P_n^{(j)} \left[\left\| x - \pi_T(x) - \sum_{i=2}^{k-1} T^{(i)} (\pi_T(x)^{\otimes i}) \right\|^2 \mathbb{I}\{B(0, h)\}(x) \right],$$

where

T : ranges in $\mathbb{G}^{p,d}$;

$T^{(i)}$: ranges in the set of i -linear maps ($2 \leq i \leq k-1$).



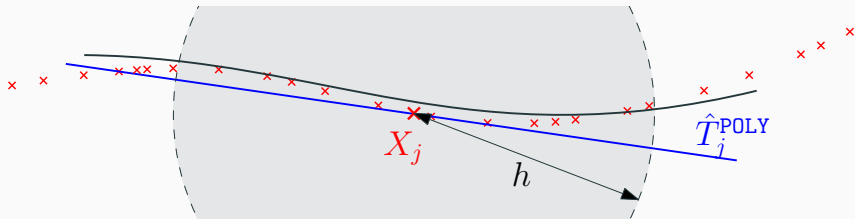
Similar methods in Cazals and Pouget 2005; Cheng and Chiu 2016; Sober and Levin 2020.

Convergence of Local Polynomials

Theorem (Aamari and Levrard 2019)

If $h = C \left(\frac{\log n}{n} \right)^{1/d}$, for all $P \in \mathcal{P}_{\text{rch}_{\min}, \mathbf{L}}^k$,

$$\mathbb{E}_{P^n} d_H(M, \hat{M}_{\text{POLY}}) \lesssim \left(\frac{\log n}{n} \right)^{\frac{k}{d}}.$$



↪ This rate is minimax optimal.

↪ Estimation of tangent spaces and curvature in the process

References

- Aamari, Eddie, Catherine Aaron, and Clément Levrard (2023). **“Minimax boundary estimation and estimation with boundary”**. In: *Bernoulli* 29.4, pp. 3334–3368.
- Aamari, Eddie and Clément Levrard (2018). **“Stability and minimax optimality of tangential Delaunay complexes for manifold reconstruction”**. In: *Discrete & computational geometry* 59.4, pp. 923–971.
- Aamari, Eddie and Clément Levrard (2019). **“Nonasymptotic rates for manifold, tangent space and curvature estimation”**. In: *Ann. statist.* 47.1, pp. 177–204. ISSN: 0090-5364. DOI: [10.1214/18-AOS1685](https://doi.org/10.1214/18-AOS1685). URL: <https://doi.org/10.1214/18-AOS1685>.
- Aizenbud, Yariv and Barak Sober (2021). **“Non-parametric estimation of manifolds from noisy data”**. In: *Arxiv preprint arxiv:2105.04754*.
- Arias-Castro, E., G. Chen, and G. Lerman (2011). **“Spectral clustering based on local linear approximations”**. In: *Electron. j. statist.* 5, pp. 1537–1587.

- Bennett, Robert (1969). **“The intrinsic dimensionality of signal collections”**. In: *Ieee transactions on information theory* 15.5, pp. 517–525.
- Block, Adam, Zeyu Jia, Yury Polyanskiy, and Alexander Rakhlin (2021). **“Intrinsic dimension estimation using wasserstein distances”**. In: *Arxiv e-prints*, arXiv–2106.
- Camastra, Francesco and Antonino Staiano (2016). **“Intrinsic dimension estimation: advances and open problems”**. In: *Information sciences* 328, pp. 26–41.
- Cazals, F. and M. Pouget (2005). **“Estimating differential quantities using polynomial fitting of osculating jets”**. In: *Comput. aided geom. design* 22.2, pp. 121–146. ISSN: 0167-8396. DOI: [10.1016/j.cagd.2004.09.004](https://doi.org/10.1016/j.cagd.2004.09.004). URL: <http://dx.doi.org/10.1016/j.cagd.2004.09.004>.
- Cheng, Siu-Wing and Man-Kwun Chiu (2016). **“Tangent estimation from point samples”**. In: *Discrete comput. geom.* 56.3, pp. 505–557. ISSN: 0179-5376. DOI: [10.1007/s00454-016-9809-z](https://doi.org/10.1007/s00454-016-9809-z). URL: <http://dx.doi.org/10.1007/s00454-016-9809-z>.
- Dedecker, Jérôme and Florence Merlevède (2019). **“Behavior of the empirical wasserstein distance in $\mathbb{R}^d$ under moment conditions”**. In: *Electronic journal of probability* 24, pp. 1–32.
- Divol, Vincent (2021). **“Minimax adaptive estimation in manifold inference”**. In: *Electronic journal of statistics* 15.2, pp. 5888–5932.
- Fefferman, Charles, Sergei Ivanov, Matti Lassas, and Hariharan Narayanan (2023). **“Fitting a manifold of large reach to noisy data”**. In: *Journal of topology and analysis*, pp. 1–82.

- Fukunaga, Keinosuke and David R Olsen (1971). **“An algorithm for finding intrinsic dimensionality of data”**. In: *Ieee transactions on computers* 100.2, pp. 176–183.
- Genovese, Christopher R, Marco Perone Pacifico, Isabella Verdinelli, Larry Wasserman, et al. (2012a). **“Minimax manifold estimation”**. In: *Journal of machine learning research* 13, pp. 1263–1291.
- Genovese, Christopher R, Marco Perone-Pacifico, Isabella Verdinelli, and Larry Wasserman (2012b). **“Manifold estimation and singular deconvolution under hausdorff loss”**. In: *The annals of statistics* 40.2, pp. 941–963.
- Hastie, Trevor and Werner Stuetzle (1989). **“Principal curves”**. In: *Journal of the american statistical association* 84.406, pp. 502–516.
- Kégl, Balázs (2002). **“Intrinsic dimension estimation using packing numbers”**. In: *Advances in neural information processing systems* 15.
- Kim, Arlene KH and Harrison H Zhou (2015). **“Tight minimax rates for manifold estimation under hausdorff loss”**. In: .
- Kim, Jisu, Alessandro Rinaldo, and Larry Wasserman (2019). **“Minimax rates for estimating the dimension of a manifold”**. In: *Journal of computational geometry* 10.1, pp. 42–95.
- Puchkin, Nikita and Vladimir G Spokoiny (2022). **“Structure-adaptive manifold estimation.”**. In: *J. mach. learn. res.* 23, pp. 40–1.

- Sober, Barak and David Levin (2020). **“Manifold approximation by moving least-squares projection (mmls)”**. In: *Constructive approximation* 52.3, pp. 433–478.
- Weed, Jonathan and Francis Bach (2019). **“Sharp asymptotic and finite-sample rates of convergence of empirical measures in wasserstein distance”**. In: *Bernoulli* 25.4A, pp. 2620–2648.