

Robust and Geometry-Aware Distribution Estimation via Optimal Transport

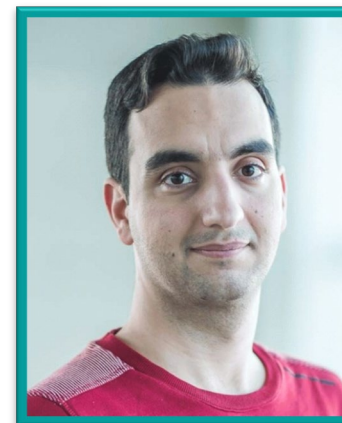
Ziv Goldfeld
Cornell University



Sloan Nietert



Rachel Cummings



Soroosh Shafiee

Machine Learning & Geometry



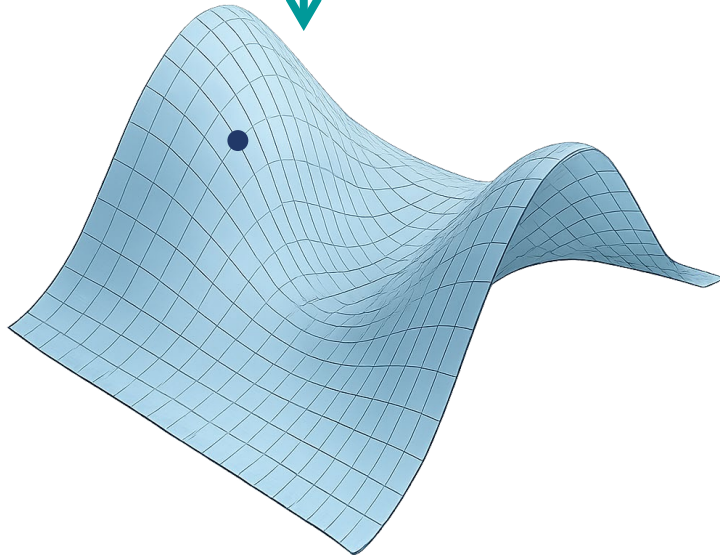
$$\begin{bmatrix} 3 \\ 1.4 \\ 1.5 \\ \vdots \\ 9.2 \end{bmatrix} \in \mathbb{R}^d$$

The quick brown fox jumps over the lazy dog

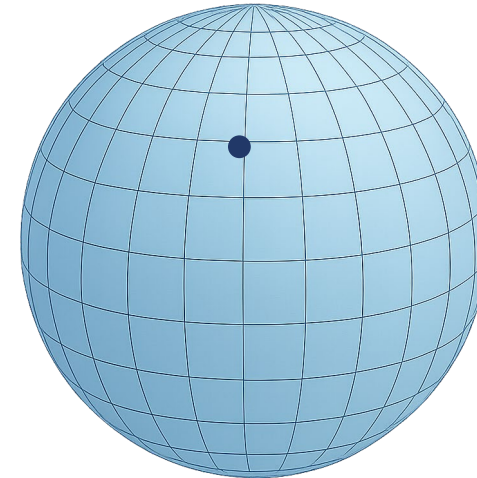


$$\begin{bmatrix} 2.7 \\ 18 \\ 2.8 \\ \vdots \\ 1.8 \end{bmatrix} \in \mathbb{R}^d$$

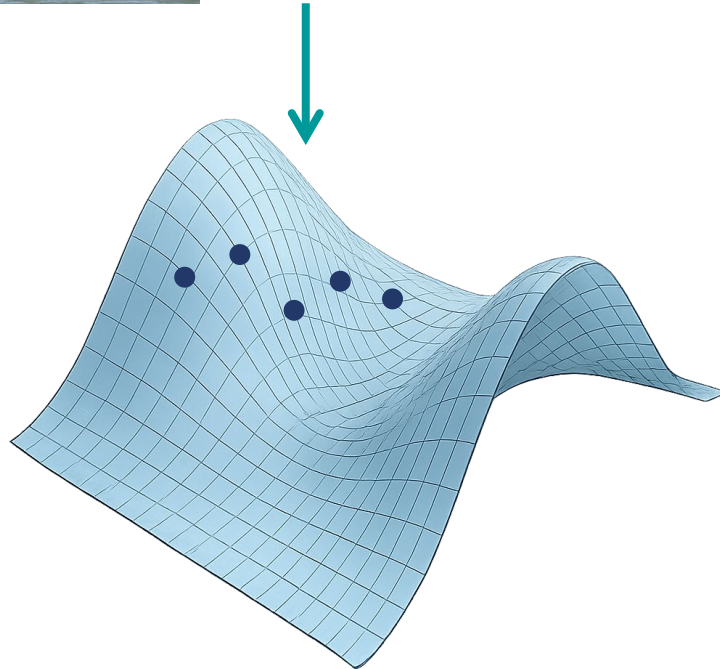
Machine Learning & Geometry



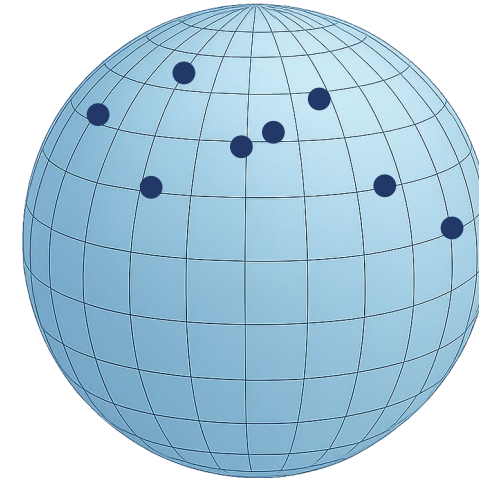
The quick brown fox jumps over the lazy dog



Machine Learning & Geometry



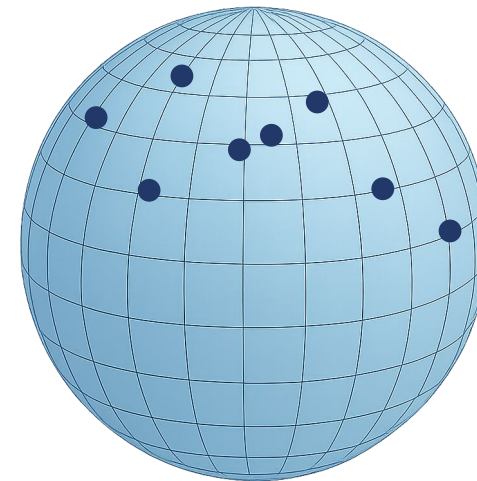
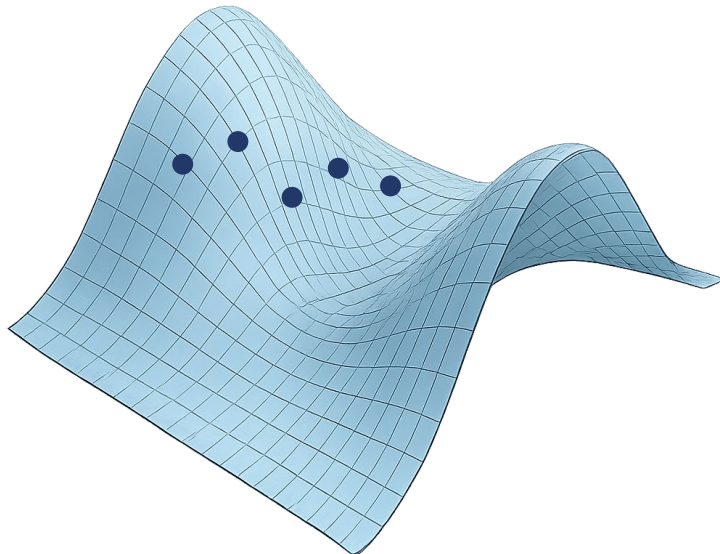
The quick brown fox jumps over the lazy dog



Machine Learning & Geometry



The quick brown fox jumps over the lazy dog



Machine Learning & Geometry



μ

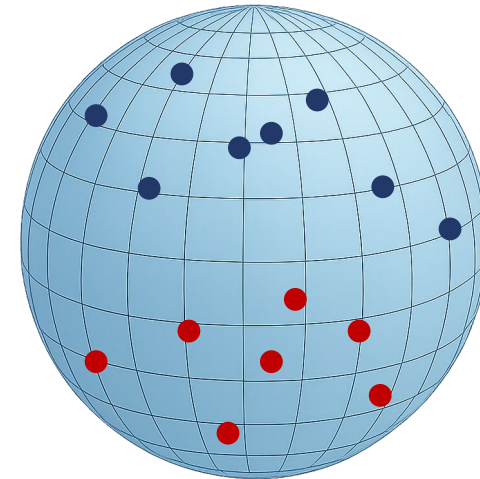
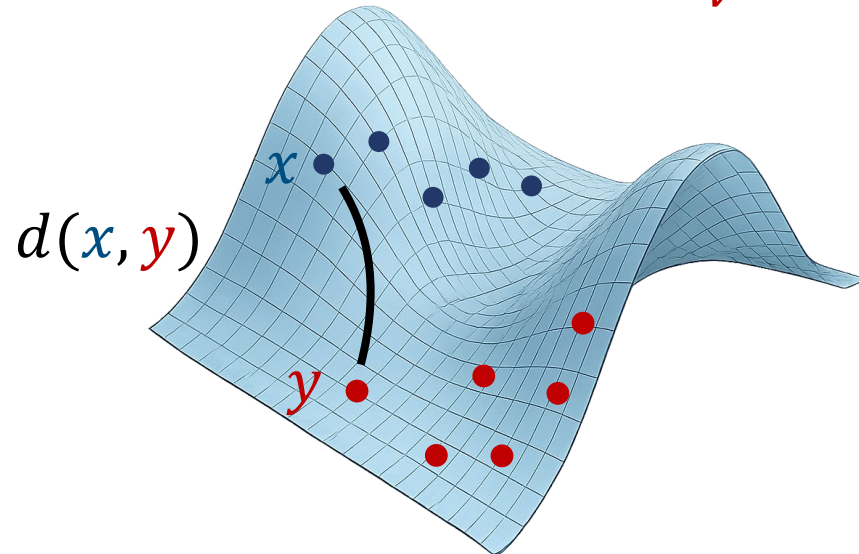
$D(\mu, \nu)?$



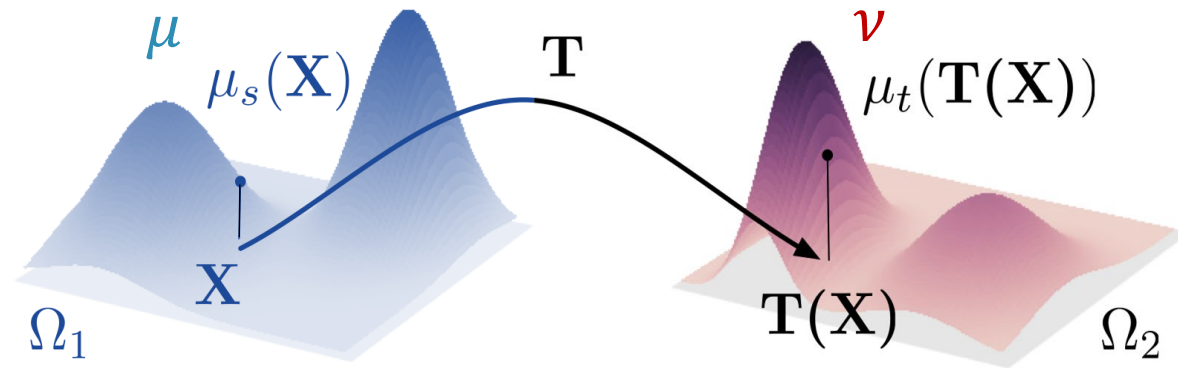
ν

The quick brown fox jumps over the lazy dog

If we are eventually to understand the capability of higher organisms for perceptual recognition...



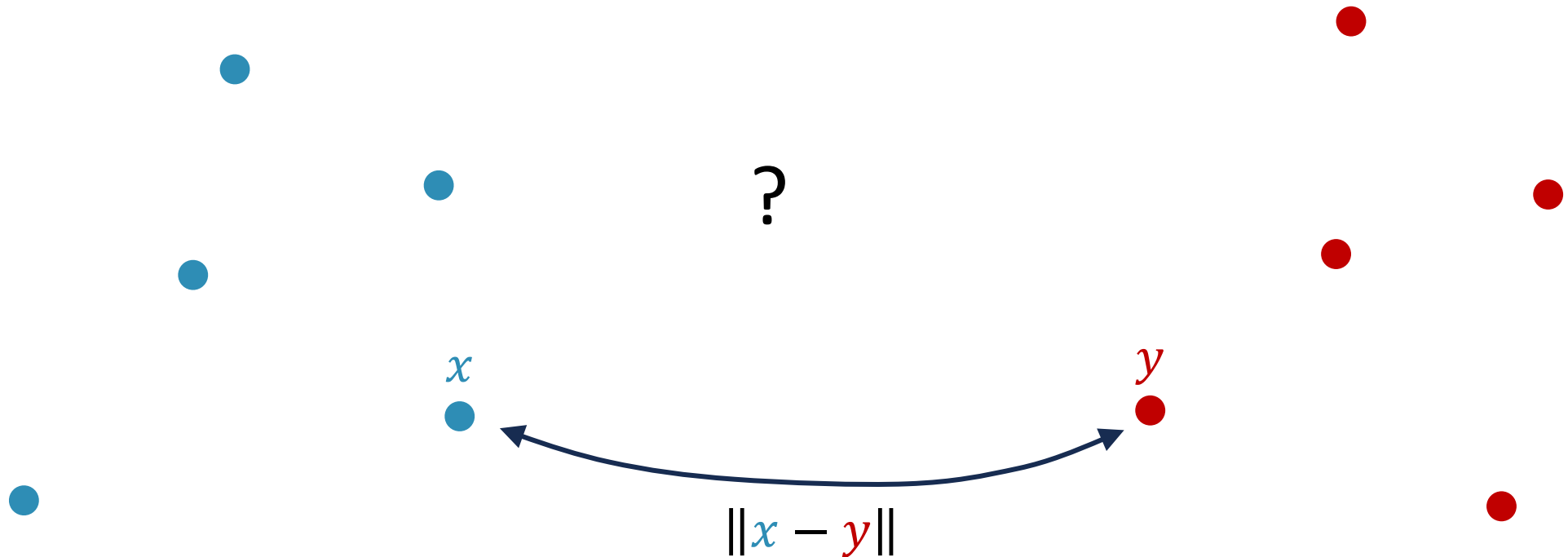
Optimal Transport



Optimal Transport

Given two dataset in $\mathbb{R}^d$, how to compare?

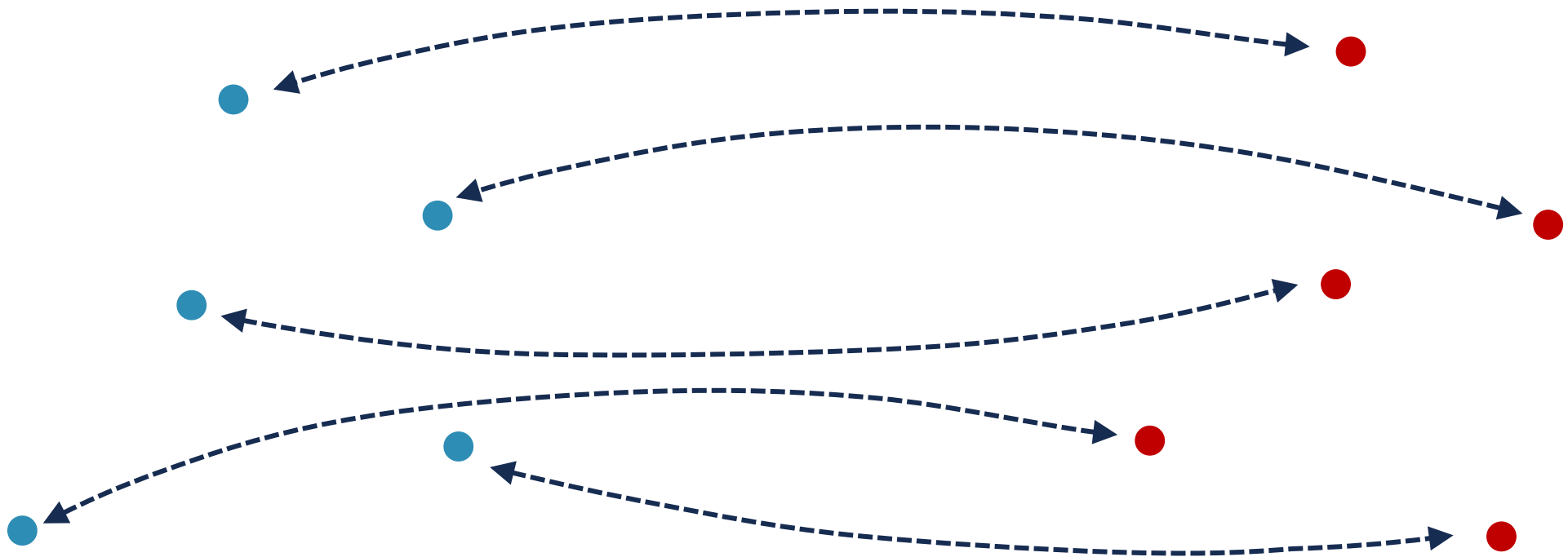
OT approach: Compute min-cost matching with distance-based cost



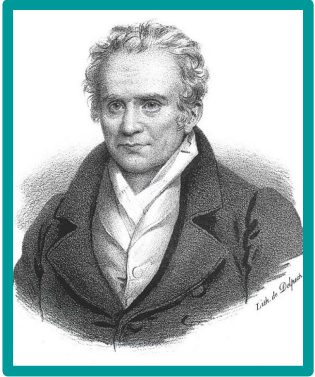
Optimal Transport

Given two dataset in $\mathbb{R}^d$, how to compare?

OT approach: Compute min-cost matching with distance-based cost



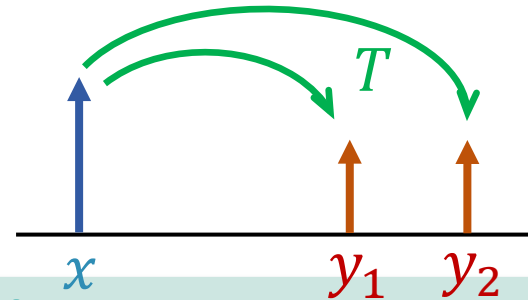
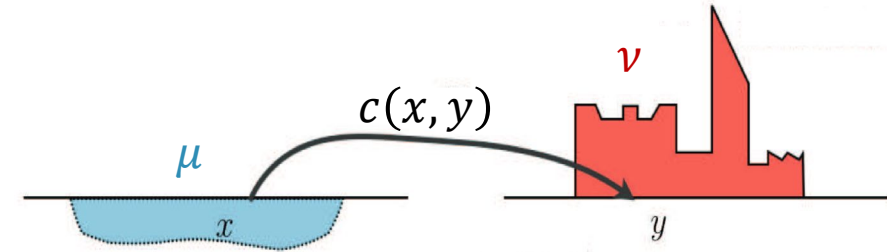
Monge & Kantorovich Formulations



Monge (1781)

$$M_c(\mu, \nu) := \inf_{\substack{T: \mathcal{X} \rightarrow \mathcal{Y} \\ T_{\#}\mu = \nu}} \mathbb{E}_{X \sim \mu} [c(X, T(X))]$$

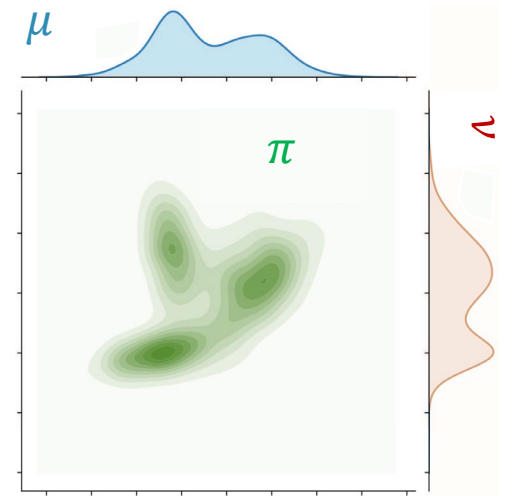
Transport map:
 $\text{Law}(T(X)) = \nu$



Kantorovich Optimal Transport

$$OT_c(\mu, \nu) := \inf_{\pi \in \Pi(\mu, \nu)} \mathbb{E}_{(X, Y) \sim \pi} [c(X, Y)]$$

Coupling (transport plan):
 $\pi(x, y) = \mu(x)\pi(y|x)$



Kantorovich (1942)

The Wasserstein Metric

p -Wasserstein Distance ($p \in [1, \infty)$)

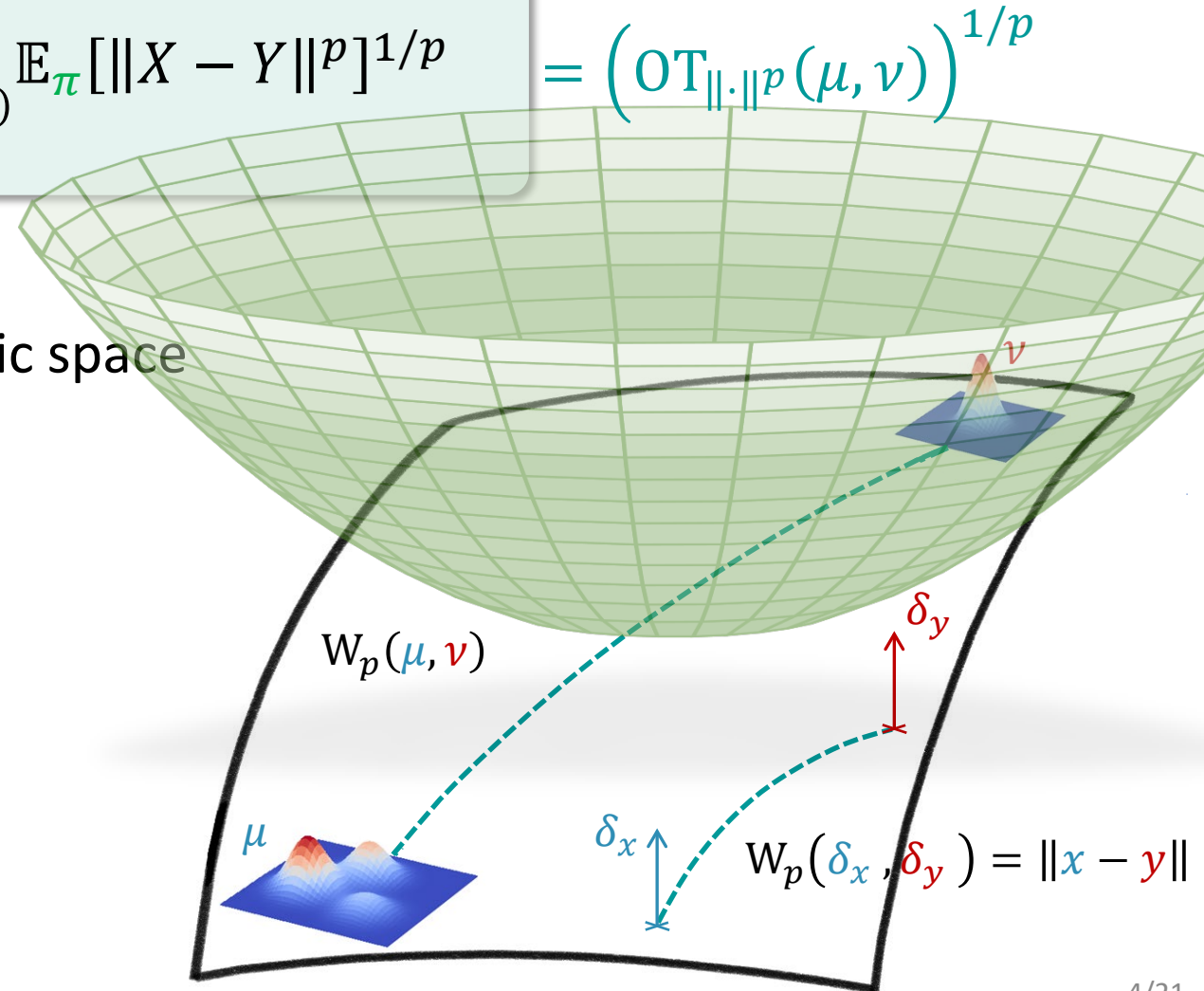
$$W_p(\mu, \nu) := \inf_{\pi \in \Pi(\mu, \nu)} \mathbb{E}_{\pi}[\|X - Y\|^p]^{1/p} = \left(\text{OT}_{\|\cdot\|^p}(\mu, \nu) \right)^{1/p}$$

Wasserstein space: $\mathfrak{W}_p = (\mathcal{P}_p(\mathbb{R}^d), W_p)$ metric space

Wasserstein geometry:

- Geodesics
- Barycenters
- Projections
- Gradient flows

⋮

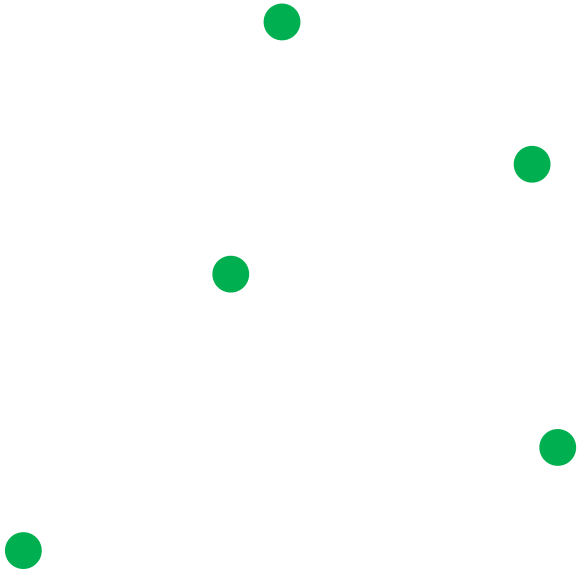


Corruption Model & Partial Optimal Transport

Corruption Model

Clean data:

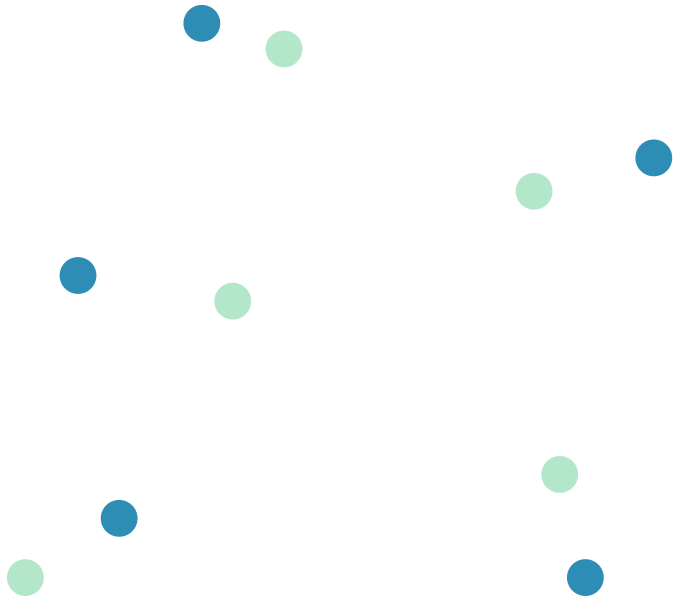
$$X_1, \dots, X_n \in \mathbb{R}^d$$



Corruption Model

Clean data: $X_1, \dots, X_n \in \mathbb{R}^d$

Local perturbation: $X'_1, \dots, X'_n \in \mathbb{R}^d$ s.t. $\frac{1}{n} \sum_{i=1}^n \|X'_i - X_i\| \leq \rho$

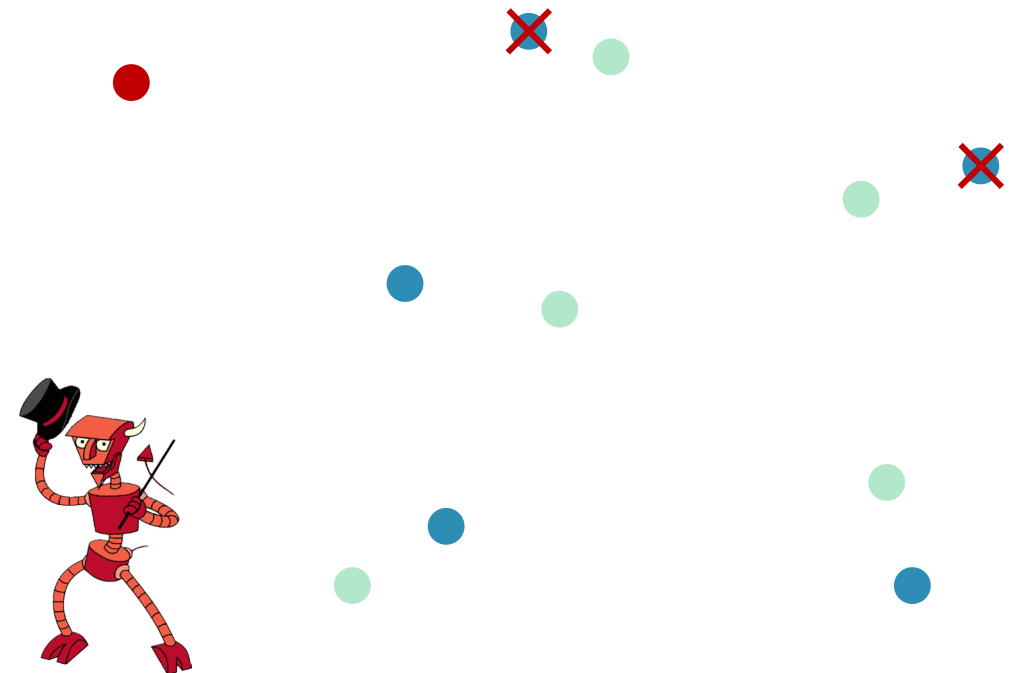


Corruption Model

Clean data: $X_1, \dots, X_n \in \mathbb{R}^d$

Local perturbation: $X'_1, \dots, X'_n \in \mathbb{R}^d$ s.t. $\frac{1}{n} \sum_{i=1}^n \|X'_i - X_i\| \leq \rho$

Global perturbation: $\tilde{X}_1, \dots, \tilde{X}_n \in \mathbb{R}^d$ s.t. $\frac{1}{n} \sum_{i=1}^n \mathbb{I}_{\{\tilde{X}_i \neq X'_i\}} \leq \epsilon$



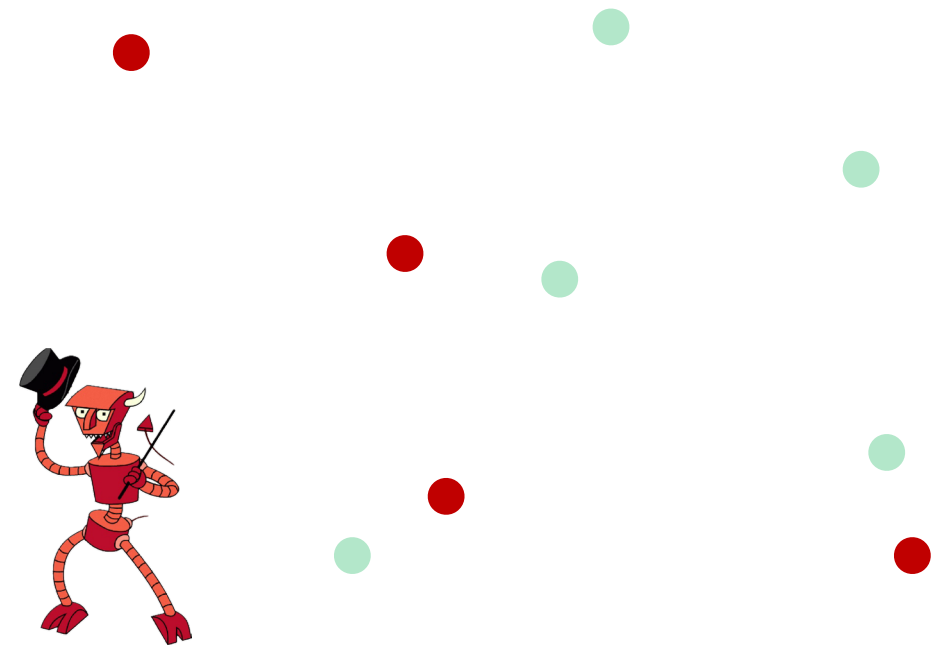
Corruption Model

Clean data: $X_1, \dots, X_n \in \mathbb{R}^d$

Local perturbation: $X'_1, \dots, X'_n \in \mathbb{R}^d$ s.t. $\frac{1}{n} \sum_{i=1}^n \|X'_i - X_i\| \leq \rho$

Global perturbation: $\tilde{X}_1, \dots, \tilde{X}_n \in \mathbb{R}^d$ s.t. $\frac{1}{n} \sum_{i=1}^n \mathbb{I}_{\{\tilde{X}_i \neq X'_i\}} \leq \epsilon$

$$\begin{aligned} \mu &= \frac{1}{n} \sum_{i=1}^n \delta_{X_i} \\ \mu' &= \frac{1}{n} \sum_{i=1}^n \delta_{X'_i} \\ \tilde{\mu} &= \frac{1}{n} \sum_{i=1}^n \delta_{\tilde{X}_i} \end{aligned} \quad \left. \begin{array}{l} W_1(\mu, \mu') \leq \rho \\ \text{TV}(\mu', \tilde{\mu}) \leq \epsilon \end{array} \right\} W_1^\epsilon(\mu, \tilde{\mu}) \leq \rho$$

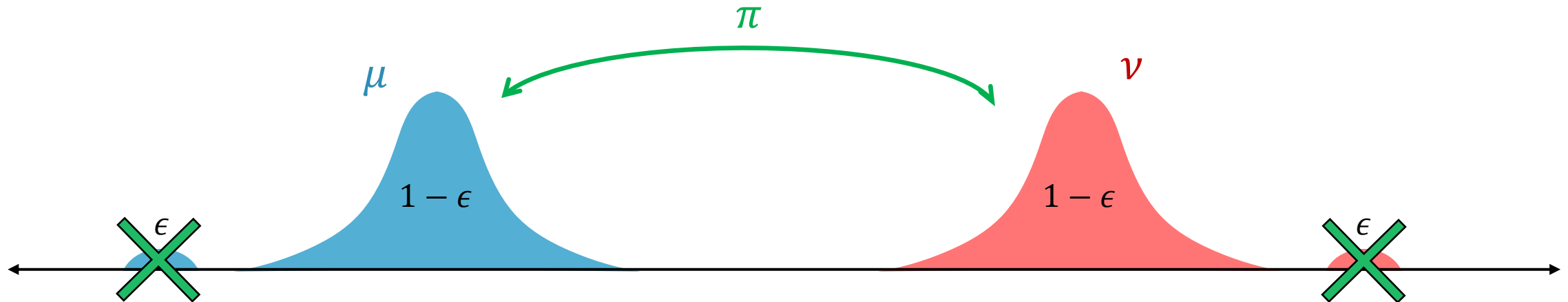


Partial Optimal Transport

Partial Optimal Transport (Caffarelli-McCann '10)

$$W_p^\epsilon(\mu, \nu) := \inf_{\pi \in \Pi_\epsilon(\mu, \nu)} \mathbb{E}_\pi[\|X - Y\|^p]^{1/p}$$

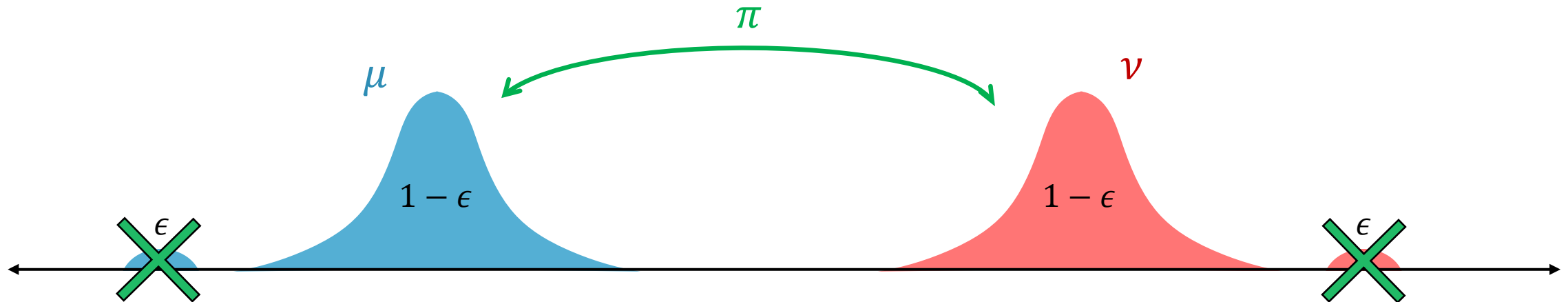
where $\Pi_\epsilon(\mu, \nu) := \{\pi \in \mathcal{M}_+(\mathcal{X} \times \mathcal{X}) : \pi_1 \leq \mu, \pi_2 \leq \nu, \pi(\mathcal{X} \times \mathcal{X}) = 1 - \epsilon\}$



Equivalent Formulations

Theorem (Nietert-Cummings-G '23)

$$W_p^\epsilon(\mu, \nu) := \inf_{\substack{\mu_-, \nu_- \in \mathcal{M}_+(\mathcal{X}): \\ \mu_- \leq \mu, \nu_- \leq \nu \\ \mu_-(\mathcal{X}) = \nu_-(\mathcal{X}) = 1 - \epsilon}} W_p(\mu_-, \nu_-) = \inf_{\substack{\mu_+, \nu_+ \in \mathcal{M}_+(\mathcal{X}): \\ \mu_+ \geq \mu, \nu_+ \geq \nu \\ \mu_+(\mathcal{X}) = \nu_+(\mathcal{X}) = 1 + \epsilon}} W_p(\mu_+, \nu_+) = \inf_{\substack{\mu' \in \mathcal{P}(\mathcal{X}): \\ \text{TV}(\mu', \mu) \leq \epsilon}} W_p(\mu', \nu)$$

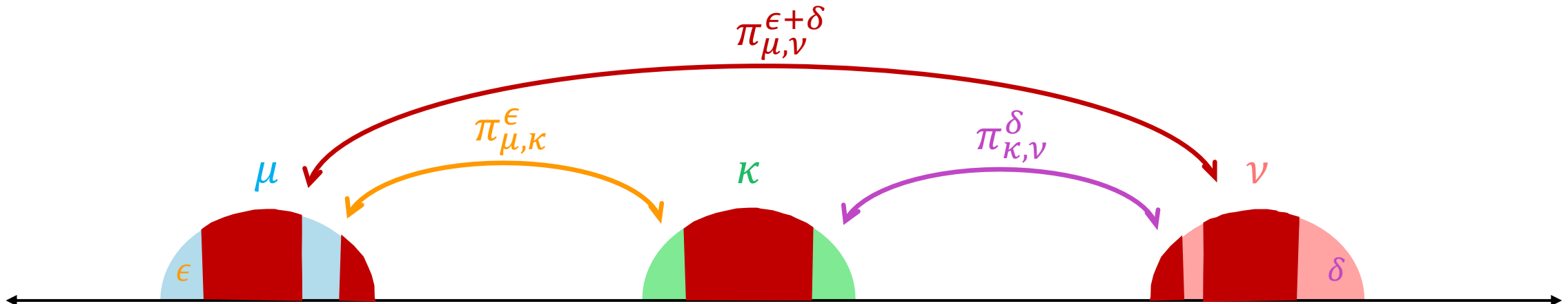


Approximate Quasi-Metric

Lemma (Nietert-Cummings-G '23)

For any $\mu, \nu, \kappa \in \mathcal{P}(\mathcal{X})$ and $\epsilon, \delta \in [0, 1]$:

- a) $W_p^\epsilon(\mu, \nu) \geq 0$ with $W_p^\epsilon(\mu, \nu) = 0$ iff $\text{TV}(\mu, \nu) \leq \epsilon$
- b) $W_p^\epsilon(\mu, \nu) = W_p^\epsilon(\nu, \mu)$
- c) $W_p^{\epsilon+\delta}(\mu, \nu) \leq W_p^\epsilon(\mu, \kappa) + W_p^\delta(\kappa, \nu)$



Existing Corruption Models

Global (TV) corruption only ($\epsilon > 0, \rho = 0$):

- Extensive statistics literature dating back to Huber ('69)
- Recent surge of interest from TCS
[Diakonikolas-Kamath-Kane-Li-Moitra-Stewart '16, '17], ...
- Closely connected to heavy-tailed & private estimation
[Hopkins-Li-Zhang '20], [Hopkins-Kamath-Majid-Narayanan '23], [Asi-Ullman-Zakynthinou '23], ...
- **Algorithmic approach:** iterative filtering / spectral reweighting

$$\hat{\mu}_{\text{IF}} = \operatorname{argmin}_{\nu: \text{TV}(\nu, \tilde{\mu}) \leq \epsilon} \lambda_{\max}(\Sigma_{\nu})$$

Existing Corruption Models

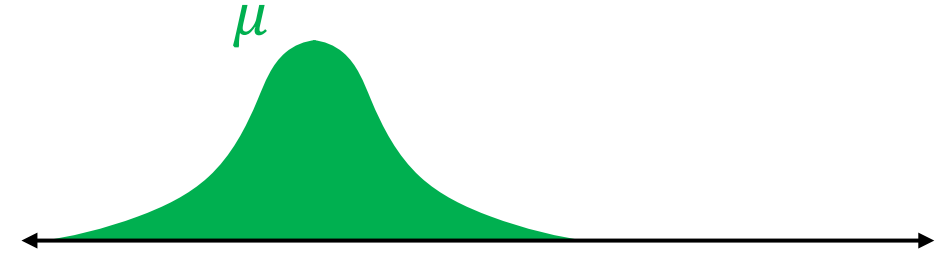
Local (Wasserstein) corruption only ($\rho > 0, \epsilon = 0$):

- Recent interest in ML from adversarial training angle
[Goodfellow-Shlens-Szegedy '15], [Madry-Makelov-Schmidt-Tsipras-Vladu '18], ...
- Explored in OR via distributionally robust stochastic optimization
[Pflug-Wozabal '07], [Shafiee-Esfahani-Kuhn '15], ...
- Examined a bit from robust statistics perspective
[Liu-Loh '21], [Zhu-Jiao-Steinhardt '22], [Chao-Dobriban '23], ...

Robust Estimation

Robust Distribution Estimation

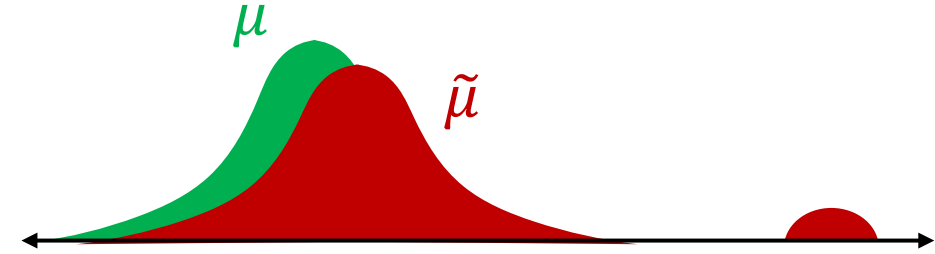
Clean distribution: $\mu \in \mathcal{G}$ (i.e., $\Sigma_\mu \preceq I_d$)



Robust Distribution Estimation

Clean distribution: $\mu \in \mathcal{G}_{\text{cov}}$ (i.e., $\Sigma_\mu \preceq I_d$)

(ρ, ϵ) -corruption: Observe $\tilde{\mu}$ s.t. $W_1^\epsilon(\mu, \tilde{\mu}) \leq \rho$



Robust estimation: Return estimate $\hat{\mu}[\tilde{\mu}] \in \mathcal{P}(\mathcal{X})$ achieving minimax risk

$$R(\rho, \epsilon, \mathcal{G}) := \inf_{\hat{\mu}} \sup_{\mu \in \mathcal{G}_{\text{cov}}} \sup_{\substack{\tilde{\mu} \in \mathcal{P}(\mathcal{X}): \\ W_1^\epsilon(\mu, \tilde{\mu}) \leq \epsilon}} W_1(\hat{\mu}[\tilde{\mu}], \mu)$$

best estimator

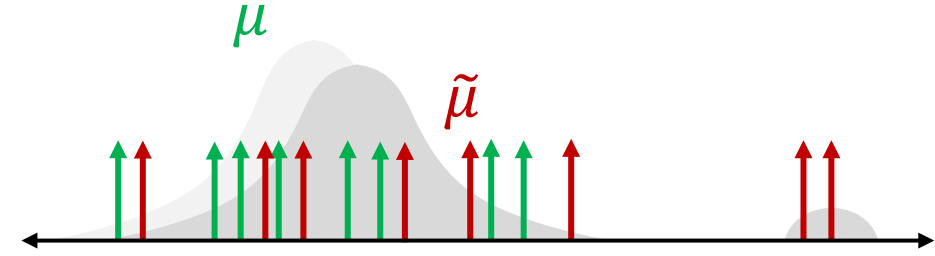
worst-case clean dist.

worst-case corruption

Robust Distribution Estimation

Clean distribution: $\mu \in \mathcal{G}_{\text{cov}}$ (i.e., $\Sigma_\mu \preceq I_d$)

(ρ, ϵ) -corruption: Observe $\tilde{\mu}$ s.t. $W_1^\epsilon(\mu, \tilde{\mu}) \leq \rho$



Robust estimation: Return estimate $\hat{\mu}[\tilde{\mu}] \in \mathcal{P}(\mathcal{X})$ achieving minimax risk

$$R(\rho, \epsilon, \mathcal{G}) := \inf_{\hat{\mu}} \sup_{\mu \in \mathcal{G}_{\text{cov}}} \sup_{\substack{\tilde{\mu} \in \mathcal{P}(\mathcal{X}): \\ W_1^\epsilon(\mu, \tilde{\mu}) \leq \rho}} W_1(\hat{\mu}[\tilde{\mu}], \mu)$$

best estimator

worst-case clean dist.

worst-case corruption

Minimax Risk Characterization

Theorem (Nietert-G.-Shafiee '24)

$$R(\rho, \epsilon, \mathcal{G}_{\text{cov}}) \asymp \sqrt{d\epsilon} + \rho$$

and there exists an efficient algorithm that achieves this risk.

Minimax Risk Characterization

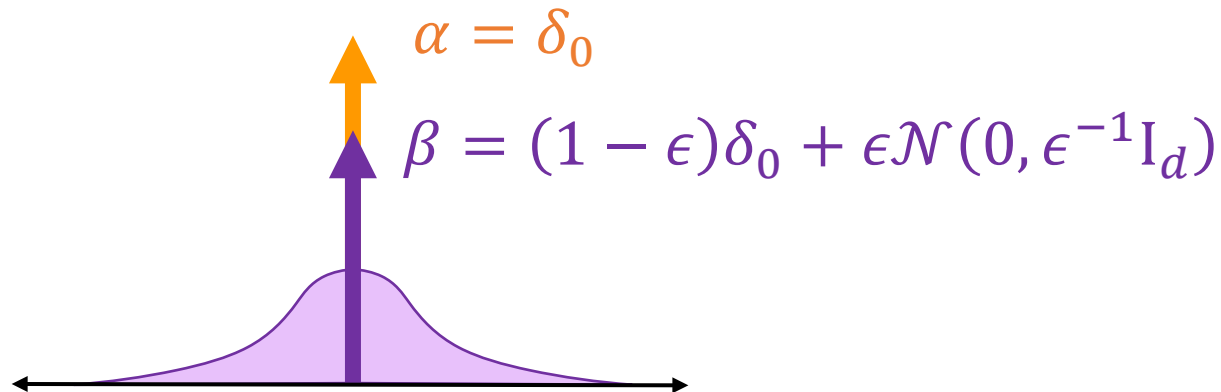
Theorem (Nietert-G.-Shafiee '24)

$$R(\rho, \epsilon, \mathcal{G}_{\text{cov}}) \asymp \sqrt{d\epsilon} + \rho$$

and there exists an efficient algorithm that achieves this risk.

$$\text{TV}(\delta_0, (1 - \epsilon)\delta_0 + \epsilon\mathcal{N}(0, \epsilon^{-1}\text{I}_d)) = \epsilon$$

$$W_1(\delta_0, (1 - \epsilon)\delta_0 + \epsilon\mathcal{N}(0, \epsilon^{-1}\text{I}_d)) = \sqrt{d\epsilon}$$



Minimax Risk Characterization

Theorem (Nietert-G.-Shafiee '24)

$$R(\rho, \epsilon, \mathcal{G}_{\text{cov}}) \asymp \sqrt{d\epsilon} + \rho$$

and there exists an efficient algorithm that achieves this risk.

$$W_1(\delta_0, \delta_\rho)$$

Minimax Risk Characterization

Theorem (Nietert-G.-Shafiee '24)

$$R(\rho, \epsilon, \mathcal{G}_{\text{cov}}) \asymp \sqrt{d\epsilon} + \rho$$

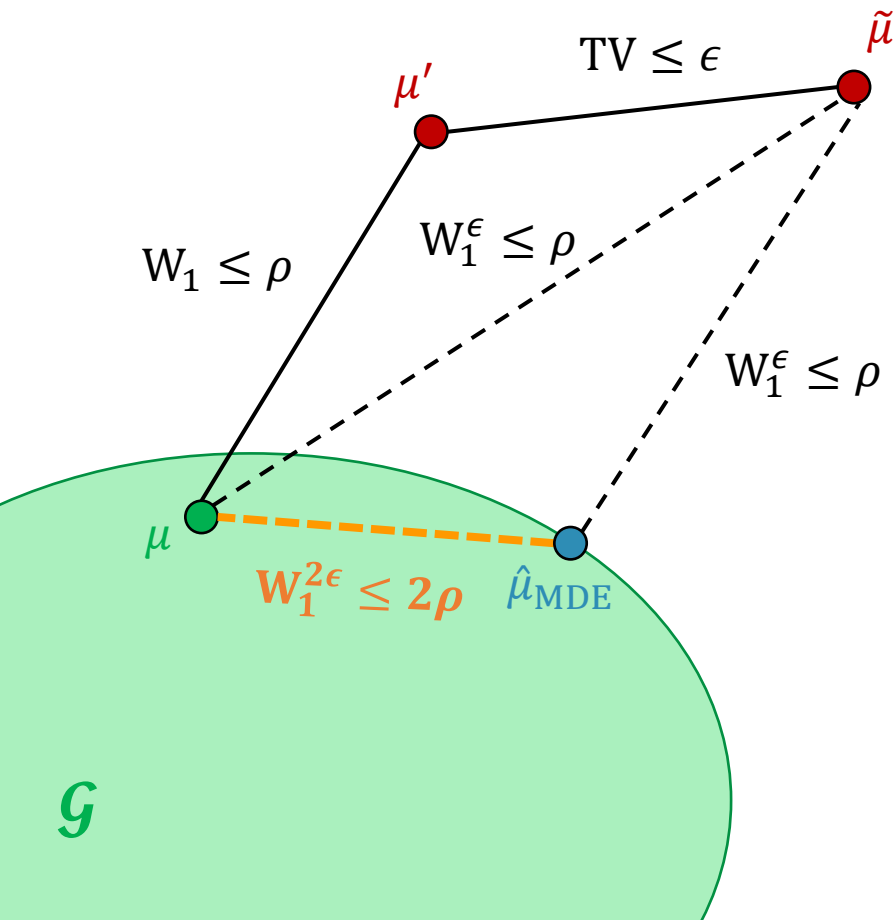
and there exists an efficient algorithm that achieves this risk.

	$\epsilon = 0$ $\rho > 0$	$\epsilon > 0$ $\rho = 0$	$\epsilon > 0$ $\rho > 0$
$\hat{\mu} = \tilde{\mu}$	$\Theta(\rho)$	∞	∞
$\hat{\mu} = \hat{\mu}_{\text{IF}}$	$O(\sqrt{d}\rho)$	$\Theta(\sqrt{d\epsilon})$	$\Theta(\sqrt{d\epsilon} + \sqrt{d}\rho)$
our work	$\Theta(\rho)$	$\Theta(\sqrt{d\epsilon})$	$\Theta(\sqrt{d\epsilon} + \rho)$

An Ideal Projection Estimator

Minimum distance estimation: Project $\tilde{\mu}$ onto $\mathcal{G}$ wrt W_1^ϵ

$$\hat{\mu}_{\text{MDE}} \in \operatorname{argmin}_{\nu \in \mathcal{G}} W_1^\epsilon(\tilde{\mu}, \nu)$$



$$W_1(\mu, \hat{\mu}_{\text{MDE}}[\tilde{\mu}]) \leq \sup_{\alpha, \beta \in \mathcal{G}: W_1^{2\epsilon}(\alpha, \beta) \leq 2\rho} W_1(\alpha, \beta)$$

$$W_1^{2\epsilon}(\alpha, \beta) \leq 2\rho$$

$$\lesssim \rho + \text{res}_{W_1}(\mathcal{G}, \epsilon)$$

$$\mathcal{G} = \mathcal{G}_{\text{cov}}$$

$$\lesssim \rho + \sqrt{d\epsilon}$$

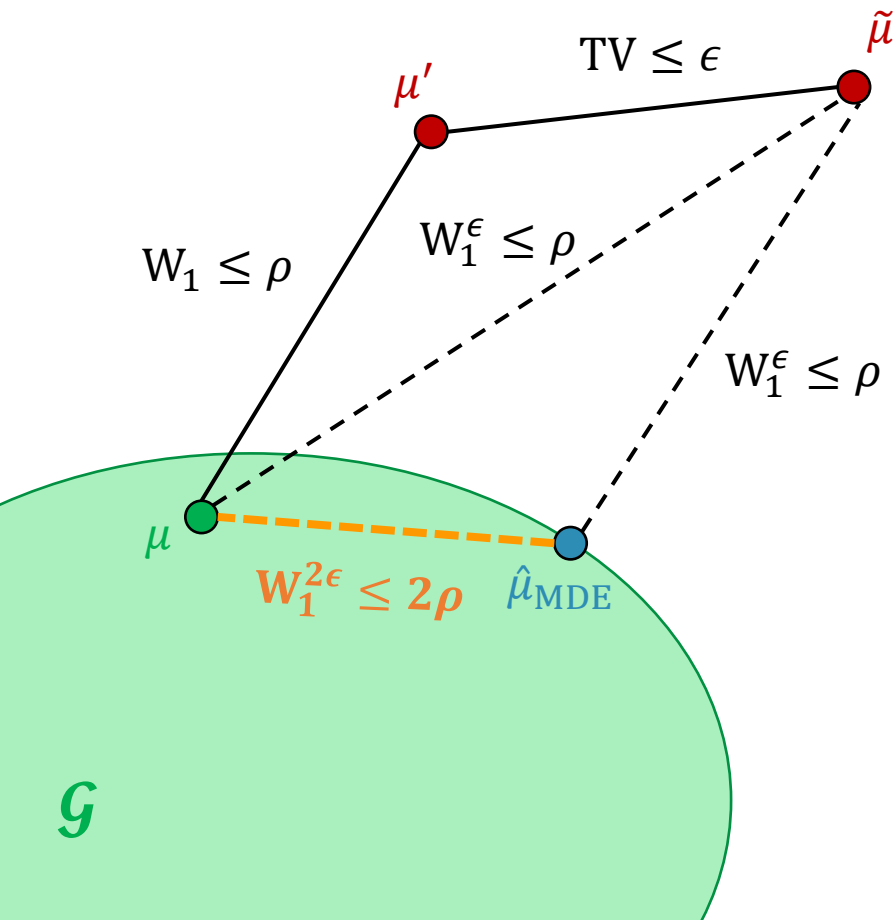
$$\text{res}_{W_1}(\mu, \epsilon) = \sup_{\mu': \mu' \leq \frac{1}{1-\epsilon}\mu} W_1(\mu, \mu')$$

An Ideal Projection Estimator

Minimum distance estimation: Project $\tilde{\mu}$ onto $\mathcal{G}$ wrt W_1^ϵ

$$\hat{\mu}_{\text{MDE}} \in \operatorname{argmin}_{\nu \in \mathcal{G}} W_1^\epsilon(\tilde{\mu}, \nu)$$

$$W_1(\hat{\mu}_{\text{MDE}}, \tilde{\mu}) \leq \rho + \sqrt{d}\delta$$



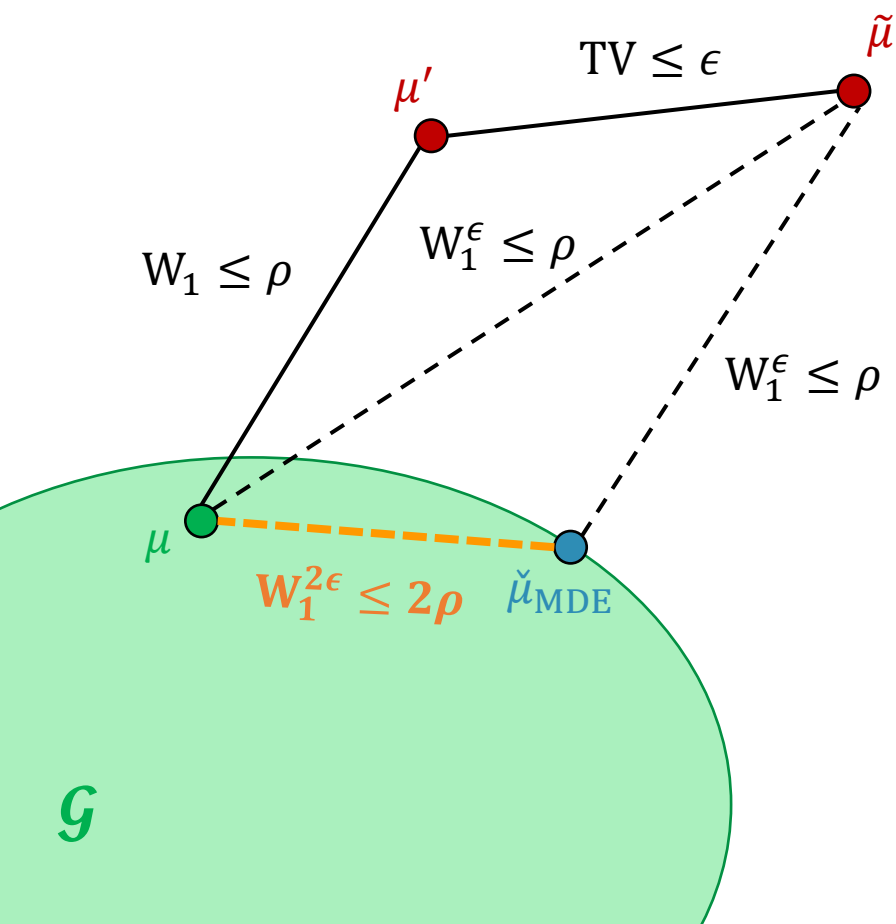
- Beyond $\mathcal{G}_{\text{cov}}$:**
- sub-Gaussianity $\delta = \tilde{O}(\epsilon)$
 - log-concavity $\delta = \tilde{O}(\epsilon)$
 - bdd. q th moments $\delta = \tilde{O}(\epsilon^{1-1/q})$

⊘ **Obstacles to efficient alg. w/ formal guarantees**

Another (tractable..?) Ideal Estimator

Use 2-Wasserstein: Project $\tilde{\mu}$ onto $\mathcal{G}$ wrt W_2^ϵ

$$\check{\mu}_{\text{MDE}} \in \operatorname{argmin}_{\nu \in \mathcal{G}} W_2^\epsilon(\tilde{\mu}, \nu)$$



- Can show also minimax optimal!

$$W_1(\check{\mu}_{\text{MDE}}, \tilde{\mu}) \leq \rho + \sqrt{d\epsilon}$$

- Key Lemma:

$$W_1(\mu, \mu') \leq \rho \implies W_2^\epsilon(\mu, \mu') \leq \frac{\rho}{\sqrt{\epsilon}}$$

Towards an Efficient Algorithm

Rewrite objective: $\min_{\nu \in \mathcal{G}_{\text{cov}}} W_2^\epsilon(\tilde{\mu}, \nu)$

$$\boxed{\min_{\nu: \text{TV}(\nu, \tilde{\mu}) \leq \epsilon} \underbrace{W_2(\nu, \mathcal{G}_{\text{cov}})}_{= \min_{\eta \in \mathcal{G}_{\text{cov}}} W_2(\nu, \eta)}} \approx \min_{\nu: \text{TV}(\nu, \tilde{\mu}) \lesssim \epsilon} \text{tr}_+(\Sigma_\nu - \text{I}_d)$$

Tractable approximation: Define $\text{tr}_+(A) := \sum_{i=1}^d \max\{\lambda_i(A), 0\}$

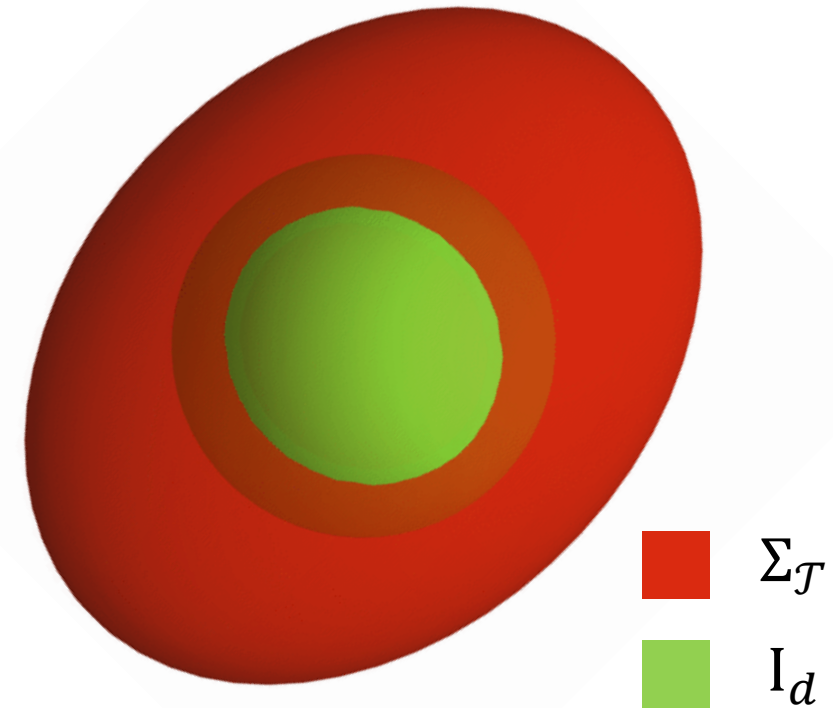
Lemma (Trace norm comparison)

$$\frac{1}{2} \text{tr}_+(\Sigma_\nu - 2\text{I}_d) \leq W_2(\nu, \mathcal{G}_{\text{cov}})^2 \leq \text{tr}_+(\Sigma_\nu - \text{I}_d)$$

The W2Project Algorithm

$$\min_{\nu: \text{TV}(\nu, \tilde{\mu}) \leq \epsilon} \text{tr}_+(\Sigma_\nu - \text{I}_d)$$

Initialize: $\mathcal{T} = \text{spt}(\tilde{\mu})$

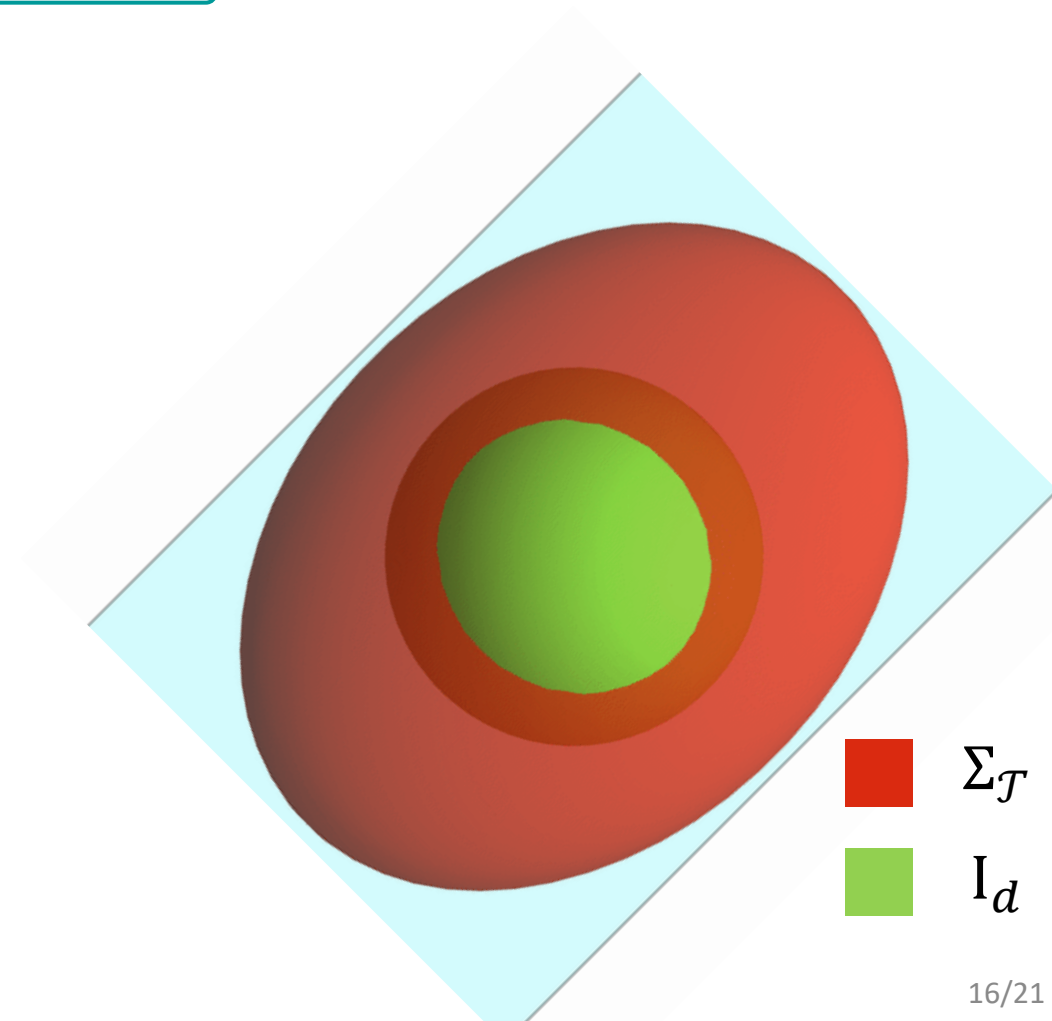


The W2Project Algorithm

$$\min_{\nu: \text{TV}(\nu, \tilde{\mu}) \leq \epsilon} \text{tr}_+(\Sigma_\nu - \text{I}_d)$$

Initialize: $\mathcal{T} = \text{spt}(\tilde{\mu})$

1. $\Pi \leftarrow$ projection onto positive eigenspace of $\Sigma_\nu - \text{I}_d$
2. Set $g(x) = \|\Pi(x - \text{mean}(\mathcal{T}))\|$ for $x \in \mathcal{T}$

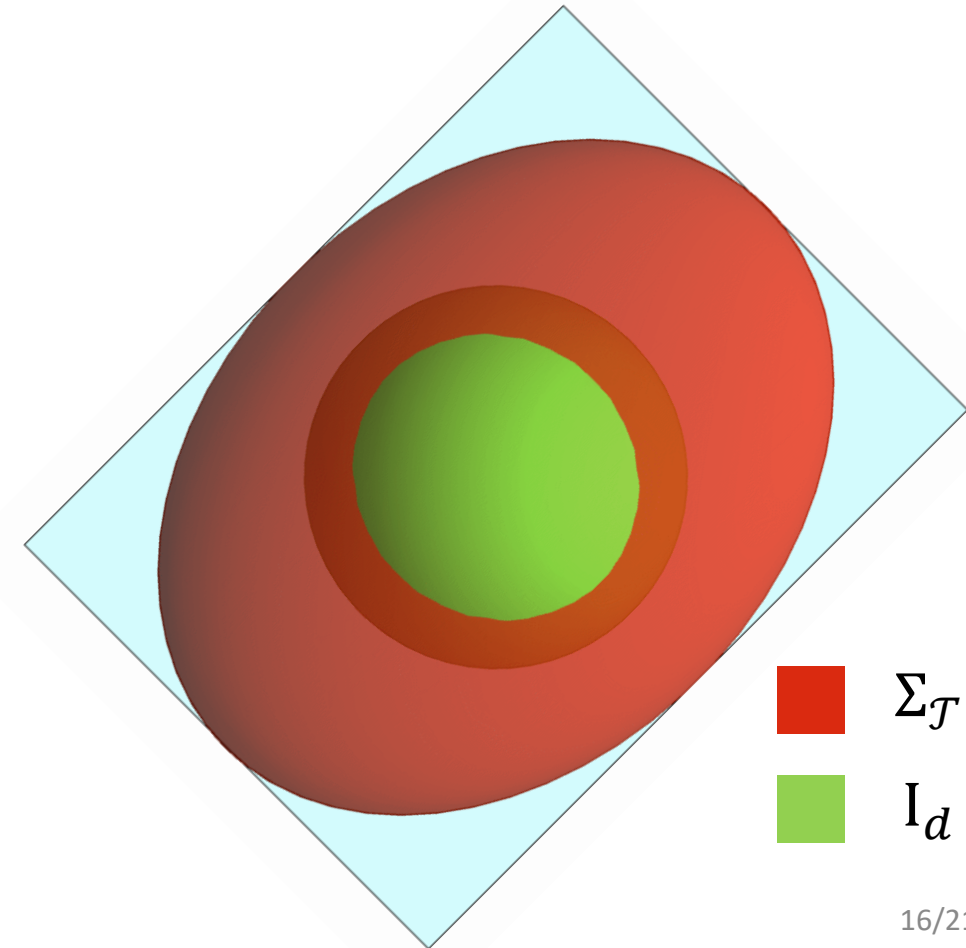


The W2Project Algorithm

$$\min_{\nu: \text{TV}(\nu, \tilde{\mu}) \leq \epsilon} \text{tr}_+(\Sigma_\nu - \text{I}_d)$$

Initialize: $\mathcal{T} = \text{spt}(\tilde{\mu})$

1. $\Pi \leftarrow$ projection onto positive eigenspace of $\Sigma_\nu - \text{I}_d$
2. Set $g(x) = \|\Pi(x - \text{mean}(\mathcal{T}))\|$ for $x \in \mathcal{T}$
3. (Randomly) throw out points $x \in \mathcal{T}$ s.t. $g(x)$ is large

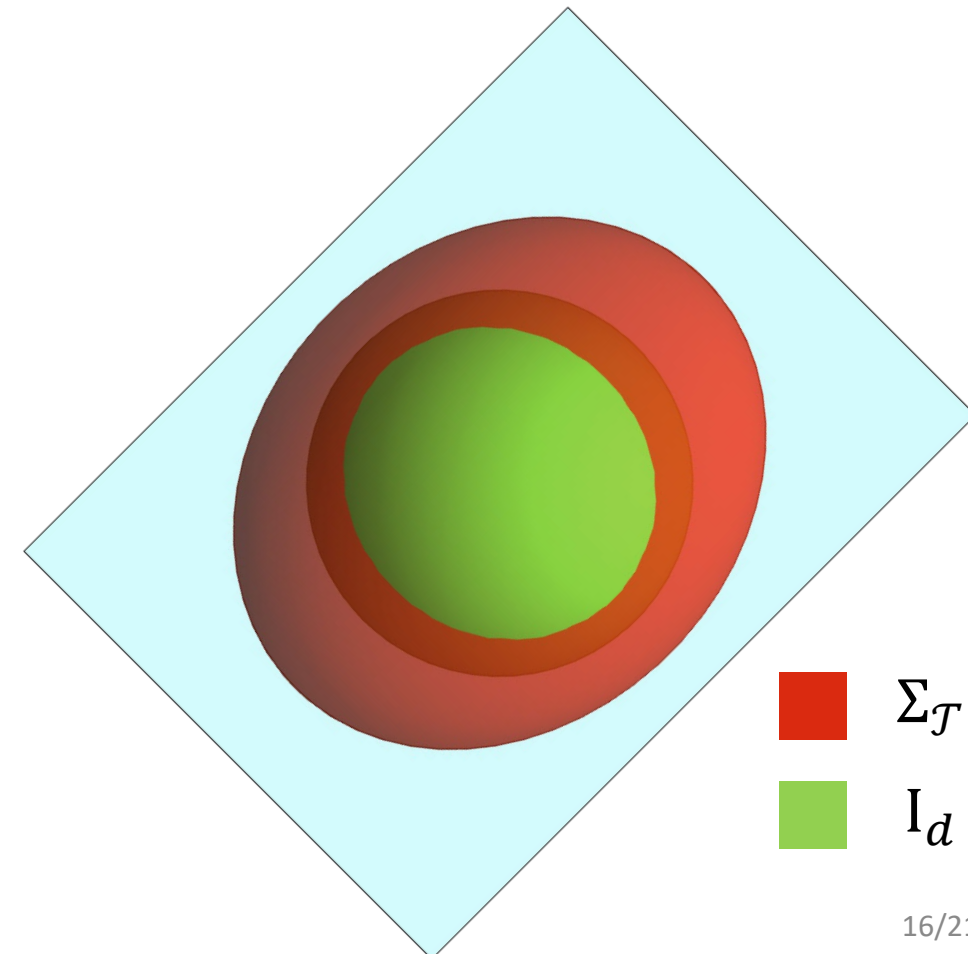


The W2Project Algorithm

$$\min_{\nu: \text{TV}(\nu, \tilde{\mu}) \leq \epsilon} \text{tr}_+(\Sigma_\nu - \text{I}_d)$$

Initialize: $\mathcal{T} = \text{spt}(\tilde{\mu})$

1. $\Pi \leftarrow$ projection onto positive eigenspace of $\Sigma_\nu - \text{I}_d$
2. Set $g(x) = \|\Pi(x - \text{mean}(\mathcal{T}))\|$ for $x \in \mathcal{T}$
3. (Randomly) throw out points $x \in \mathcal{T}$ s.t. $g(x)$ is large
4. Repeat from 1



The W2Project Algorithm

$$\min_{\nu: \text{TV}(\nu, \tilde{\mu}) \leq \epsilon} \text{tr}_+(\Sigma_\nu - I_d)$$

Algorithm 1: W2PROJECT

Input: Contamination levels ϵ and ρ , uniform discrete measure $\tilde{P}$ supported on set $T \subseteq \mathbb{R}^d$, $|T| = n$

Output: Uniform discrete measure $\hat{P}$

```
1 Compute eigen-decomposition  $\lambda_1, \dots, \lambda_d \in \mathbb{R}$ ,  $v_1, \dots, v_d \in \mathbb{R}^d$  of  $\Sigma_T - I_d$ 
2  $\Pi \leftarrow \sum_{i: \lambda_i \geq 0} v_i v_i^\top$ 
3 if  $\text{tr}(\Pi(\Sigma_T - I_d)) < C\rho^2/\epsilon$  then return  $\hat{P}_n = \text{Unif}(T)$  // LHS equals  $\text{tr}(\Sigma_T - I_d)_+$ 
4 else
5    $g(x) \leftarrow \|\Pi(x - \mu_T)\|_2^2$  for  $x \in T$ 
6   Let  $L \subseteq T$  be set of  $6\epsilon|T|$  points for which  $g(x)$  is largest
7    $f(x) \leftarrow g(x)$  for  $x \in L$  and  $f(x) \leftarrow 0$  otherwise
8   Remove each point  $x \in T$  from  $T$  with probability  $f(x) / \max_{x \in T} f(x)$ 
9   Return to Step 1 with new set  $T$ 
10 end
```

Efficient Minimax Optimal Robust Estimator

Theorem (Nietert-G.-Shafiee '24)

Fix any μ with $\Sigma_\mu \preceq I_d$. Given (ρ, ϵ) -corrupted data $\tilde{\mu}$, the W2Project algorithm produces an estimate $\hat{\mu}$ satisfying:

$$W_1(\hat{\mu}, \mu) \lesssim \sqrt{d\epsilon} + \rho$$

Efficient Minimax Optimal Robust Estimator

Theorem (Nietert-G.-Shafiee '24)

Fix any μ with $\Sigma_\mu \preceq I_d$. Given n (ρ, ϵ) -corrupted **samples**, the W2Project algorithm produces an estimate $\hat{\mu}$ satisfying:

$$W_1(\hat{\mu}, \mu) \lesssim \sqrt{d\epsilon} + \rho + \tilde{O}(n^{-1/d\vee 2})$$

with probability at least 0.99.

Empirical convergence
rate under W_1



- Near minimax optimal up to suboptimality of plug-in (logarithmic factors in general)

Dual Formulation

Theorem (Nietert-Cummings-G '23)

$$W_1^\epsilon(\mu, \nu) = \sup_{f \in \text{Lip}_1(\mathcal{X})} \mathbb{E}_\mu[f] - \mathbb{E}_\nu[f] - 2\epsilon \|f\|_\infty$$

Proof sketch:

$$W_1^\epsilon(\mu, \nu) := \inf_{\substack{\mu_+, \nu_+ \in \mathcal{M}_+(\mathcal{X}): \\ \mu_+ \geq \mu, \nu_+ \geq \nu \\ \mu_+(\mathcal{X}) = \nu_+(\mathcal{X}) = 1 + \epsilon}} W_1(\mu_+, \nu_+)$$

$$= \inf_{\alpha, \beta \in \mathcal{P}(\mathcal{X})} \sup_{f \in \text{Lip}_1(\mathcal{X})} \int f d\mu - \int f d\nu + \epsilon \left(\int f d\alpha - \int f d\beta \right)$$

Kantorovich duality

$$= \sup_{f \in \text{Lip}_1(\mathcal{X})} \int f d\mu - \int f d\nu + \epsilon \left(\inf_{\alpha \in \mathcal{P}(\mathcal{X})} \int f d\alpha - \sup_{\beta \in \mathcal{P}(\mathcal{X})} \int f d\beta \right)$$

Sion's minimax theorem

$$= \sup_{f \in \text{Lip}_1(\mathcal{X})} \int f d\mu - \int f d\nu - 2\epsilon \|f\|_\infty$$

Heuristic Neural Implementation

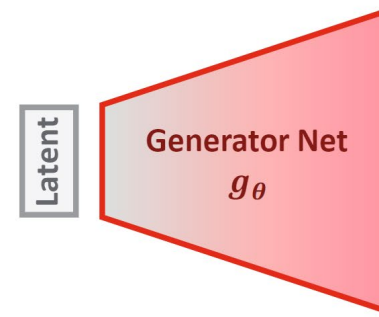
Recall: Minimax optimal (ideal) Wasserstein projection estimator

$$\inf_{\nu} W_1^{\epsilon}(\tilde{\mu}, \nu)$$

Heuristic Neural Implementation

Recall: Minimax optimal (ideal) Wasserstein projection estimator

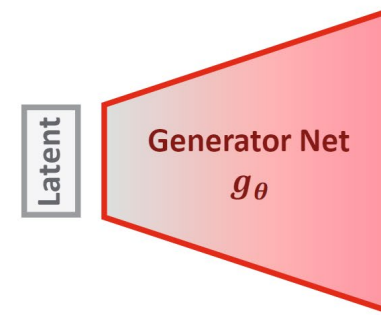
$$\inf_{\theta} W_1^{\epsilon}(\tilde{\mu}, \nu_{\theta})$$



Heuristic Neural Implementation

Recall: Minimax optimal (ideal) Wasserstein projection estimator

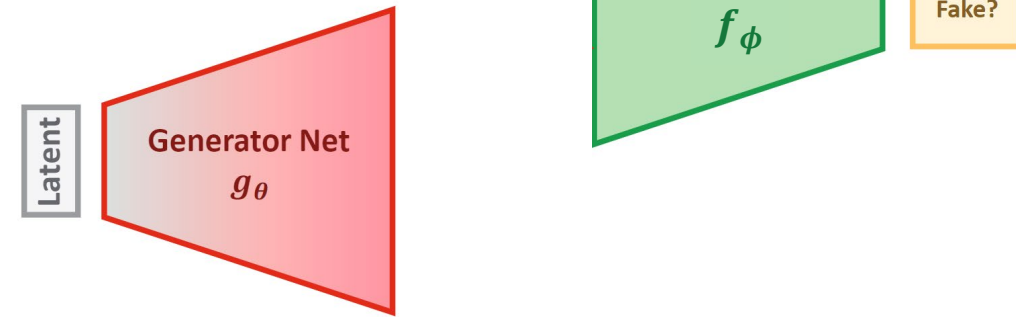
$$\inf_{\theta} W_1^{\epsilon}(\tilde{\mu}, \nu_{\theta}) \approx \inf_{\theta} \sup_{f \in \text{Lip}_1} \mathbb{E}_{\mu}[f] - \mathbb{E}_{\mathcal{N}(0, \mathbf{I})}[f \circ g_{\theta}] - 2\epsilon \|f\|_{\infty}$$



Heuristic Neural Implementation

Recall: Minimax optimal (ideal) Wasserstein projection estimator

$$\inf_{\theta} W_1^{\epsilon}(\tilde{\mu}, \nu_{\theta}) \approx \inf_{\theta} \sup_{\phi: f_{\phi} \in \text{Lip}_1} \mathbb{E}_{\mu}[f_{\phi}] - \mathbb{E}_{\mathcal{N}(0, \mathbf{I})}[f_{\phi} \circ g_{\theta}] - 2\epsilon \|f_{\phi}\|_{\infty}$$



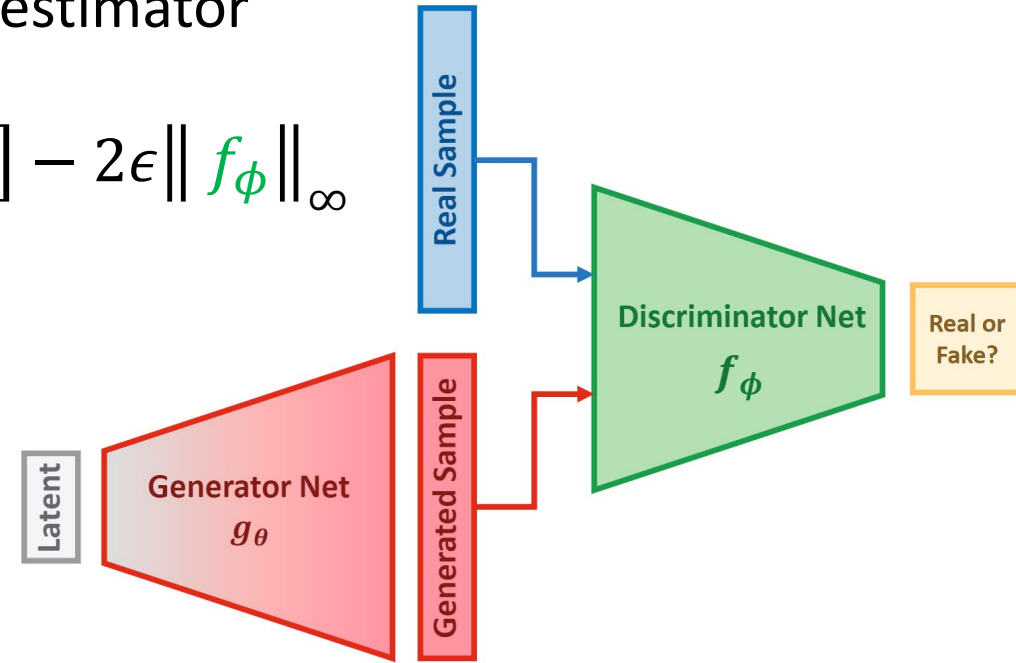
Heuristic Neural Implementation

Recall: Minimax optimal (ideal) Wasserstein projection estimator

$$\inf_{\theta} W_1^{\epsilon}(\tilde{\mu}, \nu_{\theta}) \approx \inf_{\theta} \sup_{\phi: f_{\phi} \in \text{Lip}_1} \mathbb{E}_{\mu}[f_{\phi}] - \mathbb{E}_{\mathcal{N}(0, I)}[f_{\phi} \circ g_{\theta}] - 2\epsilon \|f_{\phi}\|_{\infty}$$

⇒ Matches W-GAN form up to $\|\cdot\|_{\infty}$ penalty

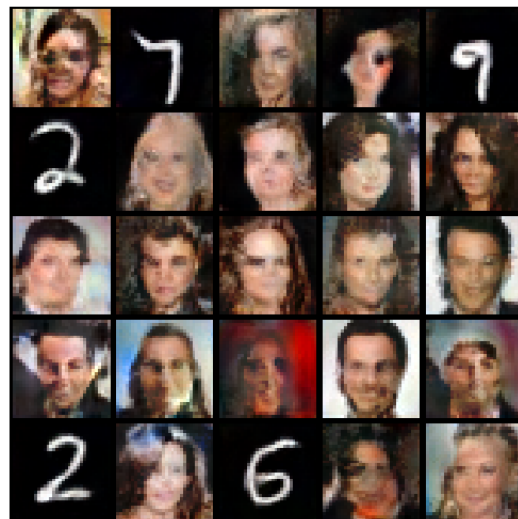
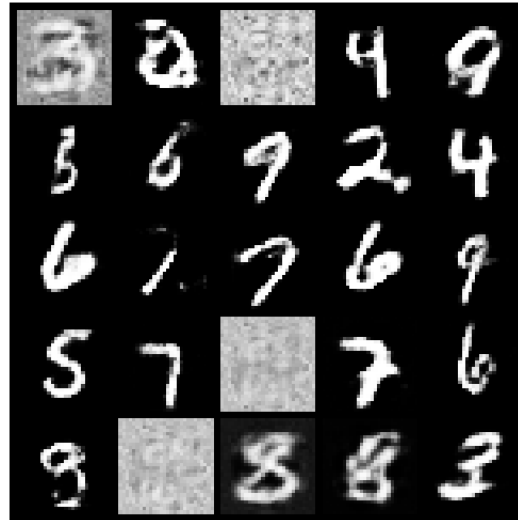
Implementation: 2 lines of code



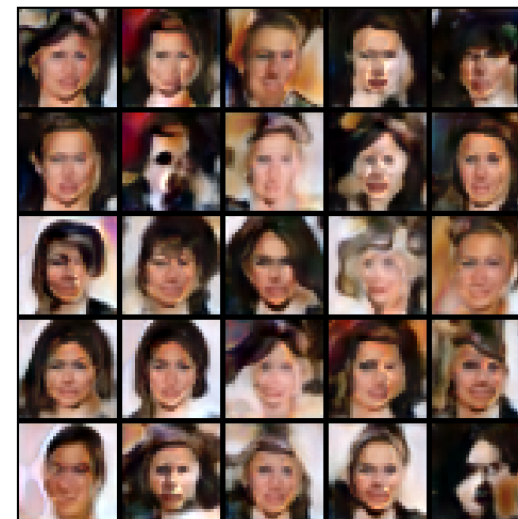
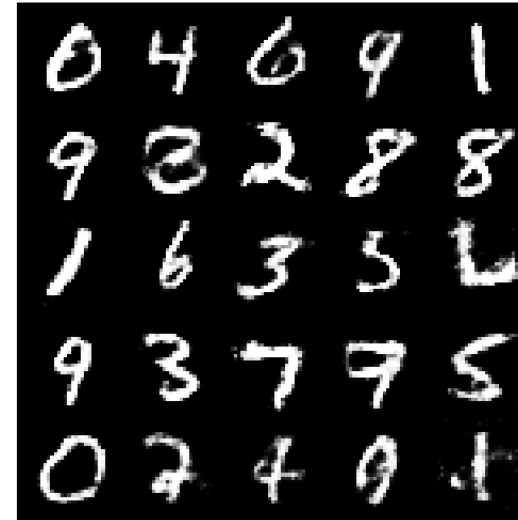
```
score = f_data.mean() - f_generated.mean() # old  
score = f_data.mean() - (1-eps)*f_generated.mean() - eps*f_data.max() # new
```

Empirical Results with TV Corruptions

Standard W-GAN



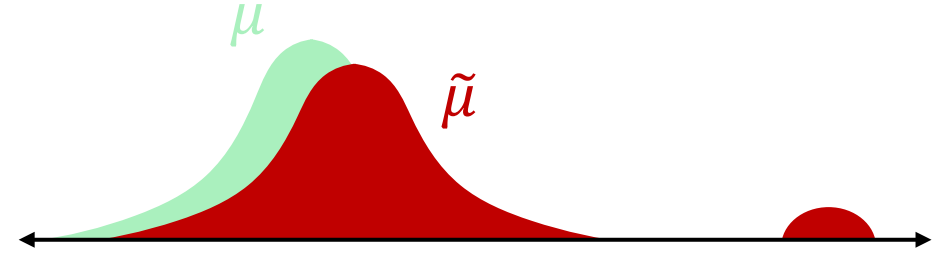
Robust W-GAN



Summary

Robust distribution estimation

- Local & global corruption model
- Under W_1 loss



Wasserstein projection estimator:

- Partial OT is compatible with filtering outliers
- W_1^ϵ or W_2^ϵ POT projection estimator is minimax optimal
- Efficient implementation via W2Project spectral algorithm

Thank you!

Neural implementation: Duality theory for POT + parametrization by neural nets

[A] S. Nietert, R. Cummings, and Z. Goldfeld, “Robust estimation under the Wasserstein distance”, ArXiv: 2302.01237

[B] S. Nietert, Z. Goldfeld, and S. Shafieezadeh-Abadeh, “Robust estimation under local and global adversarial corruptions”, COLT 2024, ArXiv:2406.06509