

Dynamic Optimal Transport

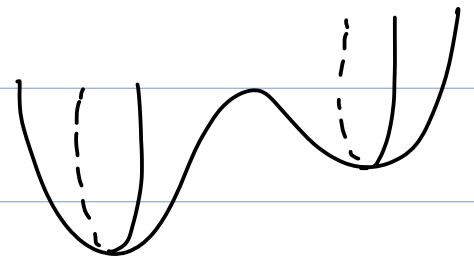
© Katy Craig, 2025

Plan

Dynamic 2-Wasserstein Metric

→ geometry

→ gradient flows



Generalizations of W_2

① Graph Wasserstein Metric

② Vector-valued Optimal Transport

References

AGS Ambrosio, Gigli, Savaré, Gradient Flows, 2005

M Mass, "Gradient flows of the entropy", 2011

Em Erbar, Maas, "Ricci Curvature...", 2012

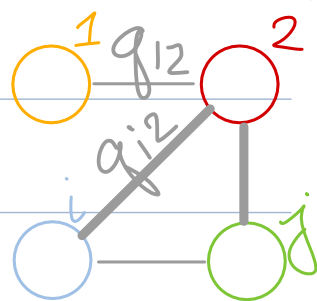
GLM Gangbo, Li, Mou, "Geodesic of Minimal...", 2017

CGN C., García Trillos, Nikolić, "Vector Valued", 2025

Recall: Graph Wasserstein Metric

Let G denote an n node, connected, symmetric, weighted graph.

$$q_{ij} = q_{ji}$$



$$\text{Let } \mathcal{P}, \tilde{\mathcal{P}} \in \mathcal{P}(G) = \left\{ \mathcal{P} \in \mathbb{R}_+^n : \sum_{i=1}^n \mathcal{P}_i = 1 \right\}$$

$$W_G^2(\mathcal{P}, \tilde{\mathcal{P}}) = \inf_{(\mathcal{P}_t, \mathcal{V}_t)} \int_0^1 \underbrace{\frac{1}{2} \sum_{i,j=0}^n |\mathcal{V}_{ijt}|^2 \check{\mathcal{P}}_{ijt} q_{ij}}_{\langle \mathcal{V}_t, \mathcal{V}_t \rangle_{\mathcal{P}_t}} dt$$

$$\partial_t \mathcal{P}_t + \text{div}_G(\mathcal{V}_t \check{\mathcal{P}}_t) = 0$$

$$\mathcal{P}_0 = \mathcal{P}, \mathcal{P}_1 = \tilde{\mathcal{P}}$$

$$\text{where } \check{\mathcal{P}} = [\check{\mathcal{P}}_{ij}]_{i,j=1}^n = [\Theta(\mathcal{P}_i, \mathcal{P}_j)]_{i,j=1}^n$$

Examples of Θ :

$$\textcircled{1} \Theta(\mathcal{P}_i, \mathcal{P}_j) \in \mathbb{R}^{\frac{\mathcal{P}_i + \mathcal{P}_j}{2}}, \textcircled{2} \Theta(\mathcal{P}_i, \mathcal{P}_j) = \sqrt{\mathcal{P}_i \mathcal{P}_j}, \dots$$

Geometry:

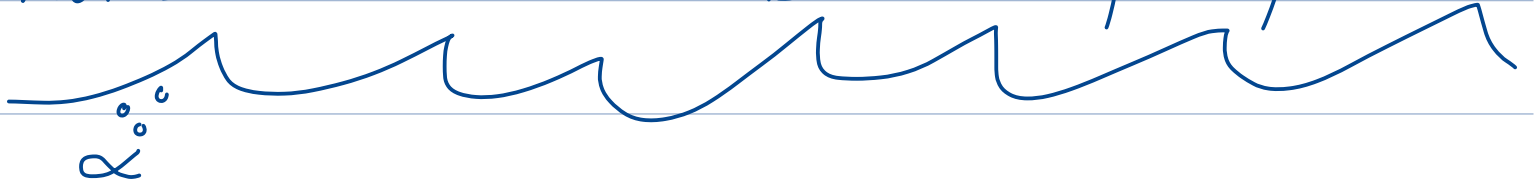
Thm (m) $(\mathcal{P}_{>0}(G), W_G)$ is a smooth Riemannian manifold.

Bad News:

- $(\mathcal{P}_{>0}(G), W_G)$ is not complete; geodesics between points in interior can touch bdy. GLM
- the minimizer of dynamic problem does not necessarily exist, but if it exists and is strictly positive, it is a geodesic. (as a curve in the space of probability measures)



For those who want a deep dive into this last statement, see pdf of previous lecture notes for more details and an open problem...



Kantorovich formulation

Thm (M) There is no Kantorovich formulation

Pf: Suppose G is a two node graph and we had a Kantorovich form.

$$\begin{aligned}\tilde{W}_G^2(p, \tilde{p}) &= \inf_{\gamma \in \Pi(p, \tilde{p})} \sum_{i,j=1}^2 d(x_i, x_j) \gamma_{ij} \\ &= \underbrace{d(x_1, x_2)}_c |p - \tilde{p}|\end{aligned}$$

Suppose $p: [0, 1] \rightarrow \mathcal{P}(G)$ is a geodesic.
Then

$$\begin{aligned}c|p_{1,t} - p_{1,s}| &\leq \tilde{W}_G^2(p_t, p_s) = |t-s|^2 \tilde{W}_G^2(p_0, p_1) \\ &= |t-s|^2 c |p_{1,0} - p_{1,1}|\end{aligned}$$

Thus $t \mapsto p_{1,t}$ is 2-Hölder cts,
hence constant; the only geodesics
in W_G are constant curves.

OTOH W_G has a rich class of
nonconstant geos $\rightarrow$ (minimizers of
dynamic form) $\square$

Gradient Flow.

$\hookrightarrow \mathbb{R}^n_+$

Given $E: P(G) \rightarrow \mathbb{R}$ smooth, $p_0 \in P(G)$
what smooth $p: [0, 1] \rightarrow P(G)$ makes
 E decrease as quickly as possible?

$$\begin{aligned} \frac{d}{dt} E(p_t) &= \left\langle \nabla_{\mathbb{R}^n} E(p_t), \partial_t p_t \right\rangle_{\mathbb{R}^n} \\ &= \left\langle \nabla_{\mathbb{R}^n} E(p_t), -\operatorname{div}_G(v_t \check{p}_t) \right\rangle_{\mathbb{R}^n} \\ &= \left\langle \nabla_G \nabla_{\mathbb{R}^n} E(p_t), v_t \check{p}_t \right\rangle_G \\ &= \left\langle \nabla_G \nabla_{\mathbb{R}^n} E(p_t), v_t \right\rangle_{p_t} \end{aligned}$$

$$\text{WGGF eqn: } \partial_t p_t - \text{div}_G (-\nabla_G \nabla_{\mathbb{R}^n} E(p_t)) p_t = 0$$

Example:

$$\Theta(p_i, p_j) = \frac{p_i - p_j}{\log p_i - \log p_j} = \int_0^1 p_i^s p_j^{1-s} ds$$

$$E(p) = \sum_i p_i \log p_i \Rightarrow \nabla_{\mathbb{R}^n} E(p) = [\log p_i + 1]_{i=1}^n$$

Given an initial condition $p_0 \in P(G)$, the WG gradient flow of E is

$$\partial_t p_t - \text{div}_G (\nabla_G [\nabla_{\mathbb{R}^n} E(p_t)] \check{p}_t) = 0$$

$$\partial_t p_t - \text{div}_G ([\log p_{i,t} - \log p_{j,t}]_{i,j=1}^n \check{p}_t) = 0$$

$$\partial_t p_t - \text{div}_G (\underbrace{[p_{i,t} - p_{j,t}]_{i,j=1}^n}_{\nabla_G p_t}) = 0$$

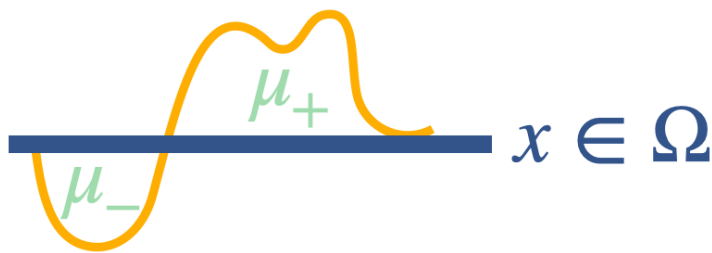
$$\partial_t p_t = \Delta_G p_t$$

Vector-Valued Optimal Transport

[C., García-Trillos, Nikolić, '25]

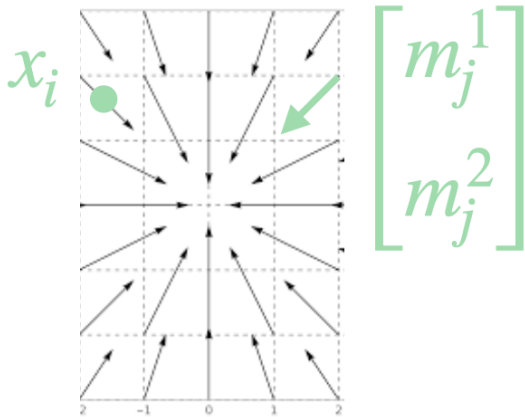
Motivation:

Signed measures



$$\mu = \begin{bmatrix} \mu_+ \\ \mu_- \end{bmatrix}$$

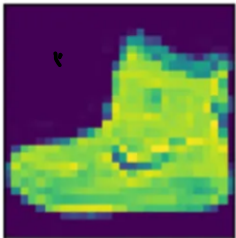
Vector fields



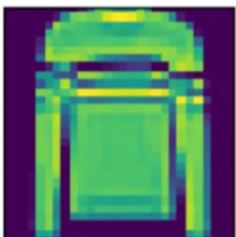
$$\mu = \sum_{i=1}^N \begin{bmatrix} m_i^1 \\ m_i^2 \end{bmatrix} \delta_{x_i}$$

Image Classification

Ankle boot



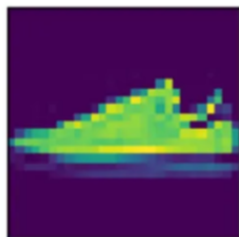
Pullover



T-shirt/top



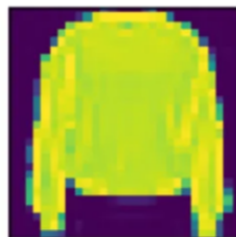
Sneaker



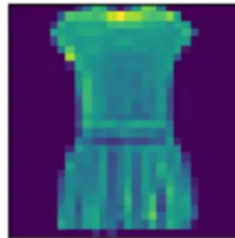
T-shirt/top



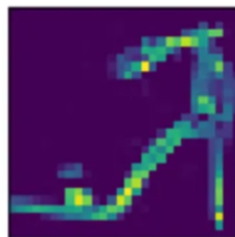
Pullover



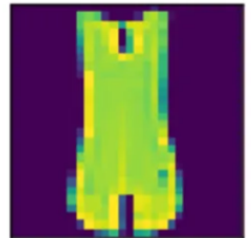
Dress



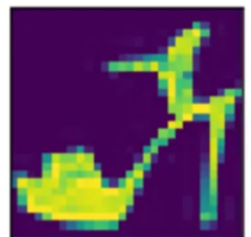
Sandal



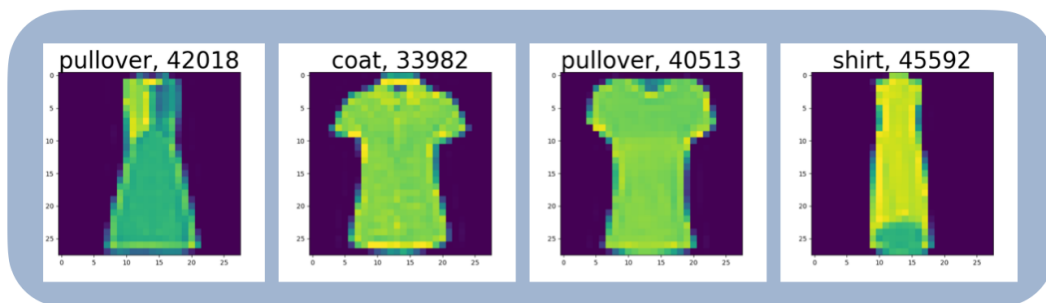
T-shirt/top



Sandal



Mild misclassifications in Fashion MNIST — for example...

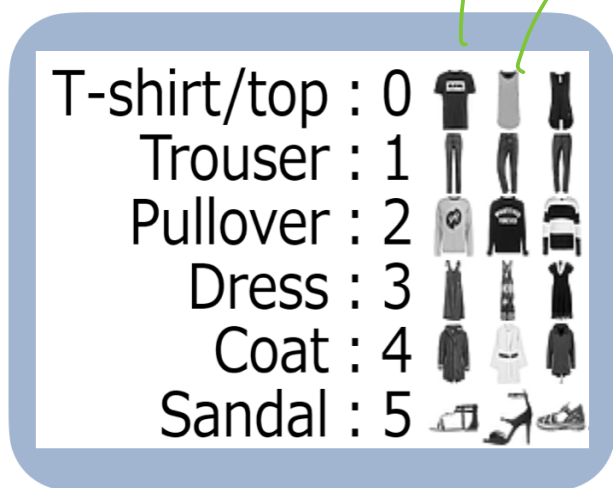


[Müller, Markert 2019]

Labeled Dataset #1



Labeled Dataset #2



Represent images as points $x_k, y_l \in \mathbb{R}^d$

Similarity between images = $\|x_k - y_l\|$

Dissimilarity between label i and label $j = q_{ij}$

Datasets are similar if their images & labels are similar.

VVV OT Formulation:

$$\mu_i = \sum_{k \in K_i} \delta_{x_k} m_k, K_i = \{k: x_k \text{ has label } i\}$$

$$\text{labeled dataset} = [\mu_i]_{i=1}^n = \underline{\mu}$$

Goal: Find a distance d on the space of vector valued measures s.t. $d([y_i]_{i=1}^n, [v_i]_{i=1}^n) \ll 1$ if the measures are similar, e.g., the datasets have similar images and labels.

Applications:

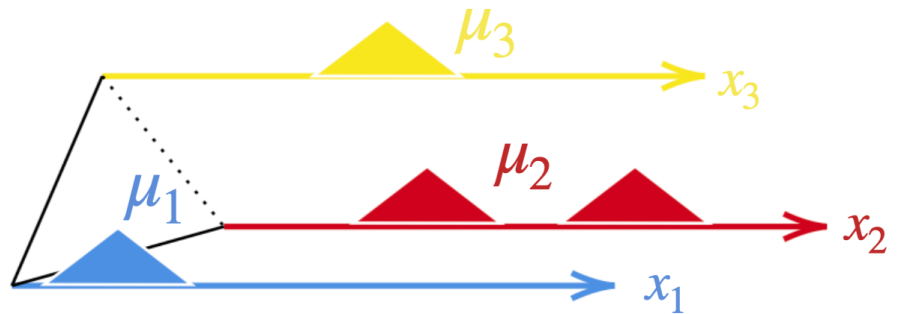
- compare the performance of two classifiers
- compare the structure of two labeled datasets
- optimization over the space of $v \mapsto v$ measures via GF.

Given $\Omega \subseteq \mathbb{R}^d$ closed, convex
 G n node sym, conn graph

We consider vv measures in

$$\mathcal{P}(\Omega \times G) = \left\{ [\mu_i]_{i=1}^n : \mu_i \in \mathcal{M}(\Omega), \sum_{i=1}^n \mu_i \in \mathcal{P}(\Omega) \right\}$$

Example:



Continuity eqn:

Given $[u_{i,t}(x)]_{i=1}^n, [v_{ij,t}(x)]_{i,j=1}^n$

$$\partial_t \rho + \nabla \cdot (\rho \underline{u}) + \text{div}_{v_G}(\check{\rho} \underline{v}) = 0$$

For $v_{ij} = -v_{ji}$, each component satisfies

$$\partial_t \rho_{i,t} + \nabla \cdot (u_{i,t} \rho_{i,t}) = \sum_{j=1}^n \theta(\rho_{i,t}, \rho_{j,t}) v_{ij,t} g_{ij}$$

Action:

$$\|(\underline{u}, \underline{v})\|_{\underline{P}}^2 := \sum_i \int_{\Omega} |\underline{u}_i|^2 p_i + \sum_{i,j} \int_{\Omega} |\underline{v}_{ij}|^2 \check{p}_{ij} q_{ij}$$

This leads to the dynamic vector valued distance...

$$W_{\Omega \times G}^2(\underline{\mu}, \underline{\nu}) = \inf_{(\underline{P}, \underline{u}, \underline{v}) \in \mathcal{C}(\underline{\mu}, \underline{\nu})} \int_0^1 \|(\underline{u}_t, \underline{v}_t)\|_{\underline{P}_t}^2 dt$$

[Thm] $W_{\Omega \times G}$ is a metric $\mathcal{P}_2(\Omega \times G)$ and a minimizer exists.

Kantorovich Formulation

- On one hand, we don't expect a Kantorovich formulation to exist, since when $\Omega = \{x_0\}$, $W_{\Omega \times G}$ reduces to WG .
- On the other hand, via analogy with Hellinger-Kantorovich, maybe there exists a static metric via liftings? nice subset of Eucl. space

Fact: $\forall \underline{\mu} \in \mathcal{P}(\Omega \times G), \exists \underline{\lambda}_{\underline{\mu}} \in \mathcal{P}(\Omega \times \mathcal{P}(G))$
s.t. $\mathbb{P} \underline{\lambda}_{\underline{\mu}} = \underline{\mu}$, where $\mathbb{P}$ is a linear projection operator.

Thm For all $\underline{\mu}, \underline{\nu} \in \mathcal{P}(\Omega \times G)$,
 $W_{\Omega \times G}(\underline{\mu}, \underline{\nu}) \leq \inf_{\substack{\mathbb{P} \underline{\lambda}_{\underline{\mu}} = \underline{\mu} \\ \mathbb{P} \underline{\lambda}_{\underline{\nu}} = \underline{\nu}}} W_{\Omega \times \mathcal{P}(G)}(\underline{\lambda}_{\underline{\mu}}, \underline{\lambda}_{\underline{\nu}})$

Good News: If Σ is bounded,

- LHS is bdd below d_{BL}
- RHS is bdd above $(d_{BL})^{1/2}$

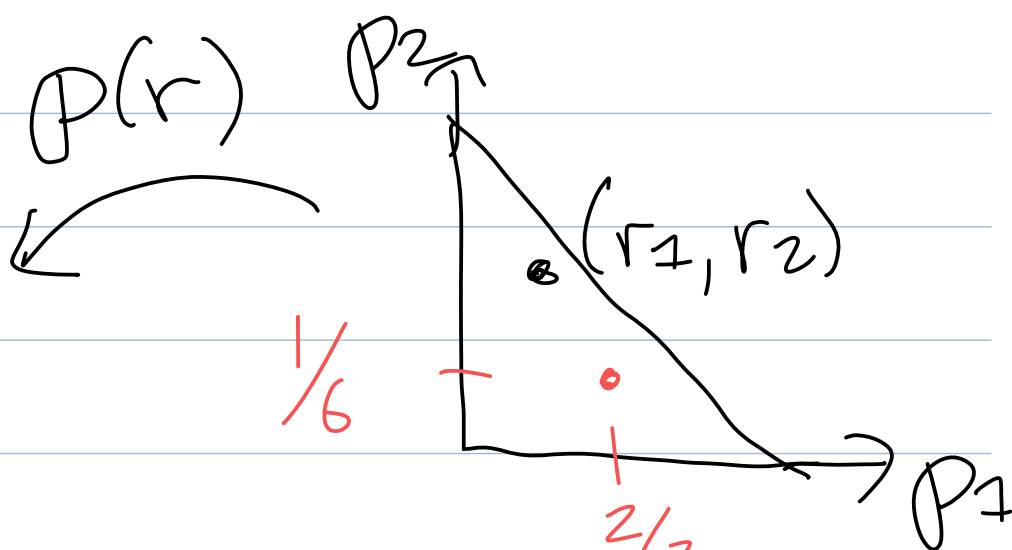
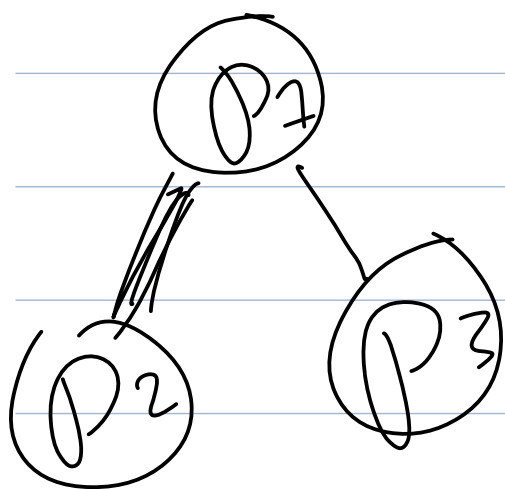
$\Rightarrow$ Topological equivalence.

... very reminiscent of relationship between W_2 and LOT

$$W_2(\mu, \nu) \leq LOT_\sigma(\mu, \nu) \leq c (W_2(\mu, \nu))^{1/6}$$

Future Work

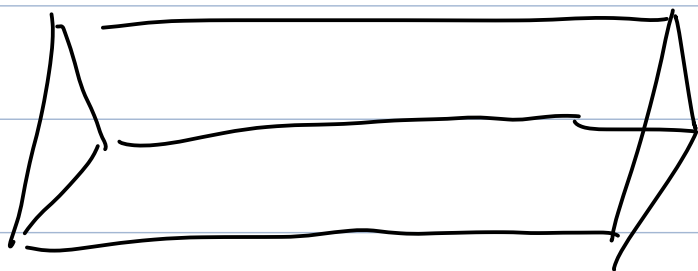
- Linearization
- Gradient flows
- Applications



$$(P(G), W_G) \cong (\Delta^{n-1}, d_{\Delta^{n-1}})$$

$$P(\Omega \times \mathcal{P}(G)) = P(\Omega \times \Delta^{n-1})$$

$\mathbb{P} \mathcal{A}$



$$(\forall \lambda)_i = d\mu_i(x) = \int_{\Delta^{n-1}} p_i(r) d\lambda(x, r)$$