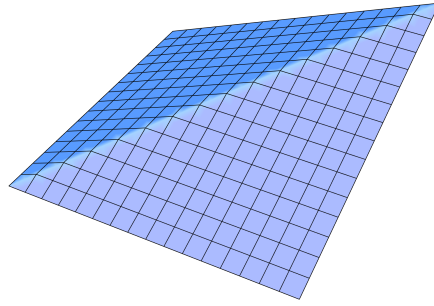




STA 4519H: Optimal Transport: Theory and Algorithms

Professor Leonard Wong • Fall 2019 • University of Toronto
 Friday, 9:00 a.m.–12:00 p.m., Sidney Smith Hall 1085

Transcribed, $\text{\LaTeX}$ ed, and edited by Rob Zimmerman

Note: These are personal lecture notes, not official course materials. They may contain errors (which by default you should attribute to me, Rob Zimmerman, rather than the course instructor). The permanent home of these — and all of my other lecture notes — is <https://rob-zimmerman.github.io/coursenotes>.

Contents

1	Foundations and Wasserstein Geometry	3
	The Monge and Kantorovich Problems	3
	Existence of Optimal Couplings	5
	Cyclical Monotonicity and c -Concavity	5
	The Fundamental Optimality Theorem	7
	Kantorovich Duality	10
	Quadratic Transport and Brenier's Theorem	11

The Wasserstein Metric	13
Dirichlet Transport and Divergences	14
Wasserstein Convergence and Displacement Interpolation	18
Dynamic Optimal Transport	21
Empirical Wasserstein Distances	23
2 Discrete Algorithms and Entropic Regularization	24
Discrete Transport as a Linear Program	24
Entropic Regularization and Sinkhorn Scaling	26
Sinkhorn Divergences and Deconvolution	29
Schrödinger Bridges and Large Deviations	30
A Dirichlet Schrödinger Problem	33
3 Riemannian and Information Geometry	35
Manifolds, Tangent Spaces, and Geodesics	35
Otto's Wasserstein Geometry	36
Fisher–Rao Geometry and Sufficiency	37
Wasserstein–Fisher–Rao Interpolation	41
Natural Gradients and Wasserstein Statistical Manifolds	44
Entropy Relaxations and Pseudo-Riemannian Connections	45
4 Stochastic Control and Extensions	46
Hamilton–Jacobi Duality	47
Stochastic Control and Girsanov's Theorem	47
Transport with Prior Dynamics	49
Schrödinger Bridges with a Markov Reference	52
Multimarginal, Unbalanced, and Martingale Transport	53
Gromov–Wasserstein Geometry	54
5 Variational Problems and Statistical Learning	56
Barycentres and Exponential Families	56
Wasserstein Barycentres	58
Metric Gradient Flows	61
Particle Discretizations	63
Minimum Kantorovich Estimation and Generative Models	64
Appendix: Lecture and Tutorial Index	68

Lecture 1 — Friday, September 13, 2019

1. Foundations and Wasserstein Geometry

The Monge and Kantorovich Problems

Let X and Y be Polish spaces with Borel probability measures $\mu \in \mathcal{P}(X)$ and $\nu \in \mathcal{P}(Y)$. In applications these measures are often empirical distributions,

$$\mu = \frac{1}{N} \sum_{i=1}^N \delta_{x_i} \quad \text{and} \quad \nu = \frac{1}{M} \sum_{j=1}^M \delta_{y_j}.$$

A **cost function** $c: X \times Y \rightarrow \mathbb{R}$ specifies the expense of moving one unit of mass from x to y . In this sense, the cost equips the probability measures with a geometry inherited from the underlying state spaces.

The original transport problem asks for a deterministic reassignment of mass.

Definition 1.1 The **Monge problem** is

$$\inf_{T: T_{\#}\mu=\nu} \int_X c(x, T(x)) \, d\mu(x),$$

where $T_{\#}\mu$ is the **pushforward measure** defined by

$$T_{\#}\mu(B) = \mu(T^{-1}(B))$$

for every Borel set $B \subseteq Y$.

The constraint in [Definition 1.1](#) is nonlinear, and a feasible map need not exist. For example, a map cannot split a point mass between two destinations. Kantorovich's relaxation replaces maps by joint distributions.

Definition 1.2 A **coupling** of μ and ν is a measure $\gamma \in \mathcal{P}(X \times Y)$ whose marginals are

μ and ν . The set of all such couplings is denoted by

$$\Pi(\mu, \nu) = \{\gamma \in \mathcal{P}(X \times Y) : (\pi_X)_\# \gamma = \mu, (\pi_Y)_\# \gamma = \nu\}.$$

The **Monge–Kantorovich problem** is

$$\mathcal{T}_c(\mu, \nu) := \inf_{\gamma \in \Pi(\mu, \nu)} \int_{X \times Y} c(x, y) \, d\gamma(x, y). \quad (1.1)$$

Disintegration writes a coupling as

$$d\gamma(x, y) = d\mu(x) \, d\gamma(y \mid x).$$

The conditional distribution $\gamma(\cdot \mid x)$ makes explicit the possibility of splitting the mass initially located at x . When $\gamma = (\text{Id}, T)_\# \mu$, this conditional distribution is the point mass $\delta_{T(x)}$, and (1.1) reduces to the Monge problem.

The distinction is illustrated in Figure 1. A **transport map** selects one destination for each source, whereas a general coupling may send fractions of the same source mass to several destinations.

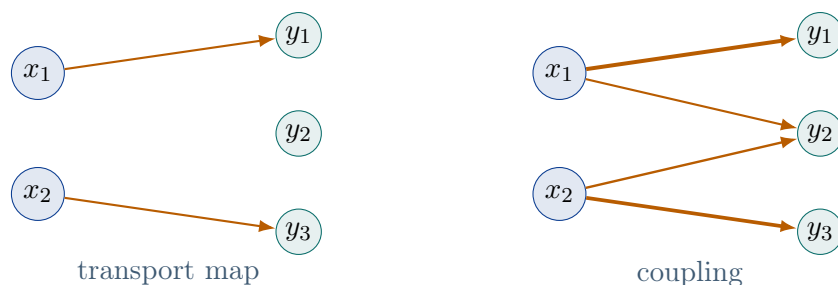


FIGURE 1

The set $\Pi(\mu, \nu)$ is convex: the marginal conditions are linear in γ . The objective in (1.1) is also linear. Thus the Kantorovich formulation is an infinite dimensional linear programming problem, and its duality theory will be central.

Important examples include

$$c(x, y) = |x - y|, \quad c(x, y) = \frac{1}{2}|x - y|^2, \quad \text{and} \quad c(x, y) = \frac{1}{2}d(x, y)^2$$

on a metric space or a Riemannian manifold.

Existence of Optimal Couplings

Theorem 1.3 Suppose that $c: X \times Y \rightarrow \mathbb{R}$ is continuous and bounded below. Then, for every $\mu \in \mathcal{P}(X)$ and $\nu \in \mathcal{P}(Y)$, the infimum in (1.1) is attained whenever its value is finite.

Proof. Choose a minimizing sequence $(\gamma_n)_{n \geq 1} \subseteq \Pi(\mu, \nu)$. The family $\Pi(\mu, \nu)$ is tight because its two marginals are fixed and tight. Prokhorov's theorem therefore gives a weakly convergent subsequence, which we continue to denote by (γ_n) , with $\gamma_n \Rightarrow \gamma$.

The marginal maps are continuous under weak convergence. Hence, for every bounded continuous f on X and g on Y ,

$$\int f(x) d\gamma = \lim_{n \rightarrow \infty} \int f(x) d\gamma_n = \int f d\mu$$

and similarly $\int g(y) d\gamma = \int g d\nu$. Thus $\gamma \in \Pi(\mu, \nu)$.

After adding a constant, we may assume that $c \geq 0$. Set $c_k = c \wedge k$. Each c_k is bounded and continuous, so

$$\int c_k d\gamma = \lim_{n \rightarrow \infty} \int c_k d\gamma_n \leq \liminf_{n \rightarrow \infty} \int c d\gamma_n.$$

Letting $k \rightarrow \infty$ and applying the monotone convergence theorem yields

$$\int c d\gamma \leq \liminf_{n \rightarrow \infty} \int c d\gamma_n = \mathcal{T}_c(\mu, \nu).$$

The reverse inequality follows from feasibility, so γ is optimal. □

Cyclical Monotonicity and c -Concavity

For the **quadratic cost** on $\mathbb{R}^n$, suppose that sending $\mathbf{x}_i$ to $\mathbf{y}_i$ is optimal. Any cyclic rearrangement of the destinations must be no cheaper, so

$$\sum_{i=1}^N \frac{1}{2} \|\mathbf{x}_i - \mathbf{y}_i\|^2 \leq \sum_{i=1}^N \frac{1}{2} \|\mathbf{x}_i - \mathbf{y}_{\sigma(i)}\|^2$$

for every cyclic permutation σ . Expanding the squares gives a monotonicity inequality. Conversely, if ϕ is convex and $\mathbf{y}_i \in \partial\phi(\mathbf{x}_i)$, then the subgradient inequalities telescope around every cycle and give precisely the same optimality condition. This motivates the following definition.

Definition 1.4 A set $\Gamma \subseteq X \times Y$ is **c -cyclically monotone** if, for every $N \geq 1$, every $(x_i, y_i)_{i=1}^N \subseteq \Gamma$, and every cyclic permutation σ of $\{1, \dots, N\}$,

$$\sum_{i=1}^N c(x_i, y_i) \leq \sum_{i=1}^N c(x_i, y_{\sigma(i)}). \quad (1.2)$$

Definition 1.5 For $\psi: Y \rightarrow \mathbb{R} \cup \{-\infty\}$, its **c -transform** is

$$\psi^c(x) = \inf_{y \in Y} \{c(x, y) - \psi(y)\}.$$

A function $\varphi: X \rightarrow \mathbb{R} \cup \{-\infty\}$ is **c -concave** if $\varphi = \psi^c$ for some ψ . Its **c -superdifferential** is

$$\partial^c \varphi = \{(x, y) \in X \times Y : \varphi(x) + \varphi^c(y) = c(x, y)\}.$$

When the fibre $\{y : (x, y) \in \partial^c \varphi\}$ is a singleton, its unique element is denoted by $\nabla^c \varphi(x)$.

The inequality $\varphi(x) + \varphi^c(y) \leq c(x, y)$ holds for all (x, y) . Equality identifies the pairs that are compatible with the potential.

For the quadratic cost on $\mathbb{R}^n$ there is a direct link with ordinary convex analysis.

Proposition 1.6 Let $c(\mathbf{x}, \mathbf{y}) = \|\mathbf{x} - \mathbf{y}\|^2/2$. A function φ is c -concave if and only if

$$\phi(\mathbf{x}) = \frac{1}{2}\|\mathbf{x}\|^2 - \varphi(\mathbf{x})$$

is lower semicontinuous and convex. Under this correspondence,

$$\partial^c \varphi = \partial \phi.$$

Proof. Expanding the cost gives

$$\varphi(\mathbf{x}) = \inf_{\mathbf{y}} \left\{ \frac{1}{2}\|\mathbf{x} - \mathbf{y}\|^2 - \psi(\mathbf{y}) \right\} = \frac{1}{2}\|\mathbf{x}\|^2 - \sup_{\mathbf{y}} \left\{ \mathbf{x} \cdot \mathbf{y} - \left(\frac{1}{2}\|\mathbf{y}\|^2 - \psi(\mathbf{y}) \right) \right\}.$$

Thus ϕ is a **convex conjugate**. Conversely, if ϕ is lower semicontinuous and convex, set

$$\psi(\mathbf{y}) = \frac{1}{2}\|\mathbf{y}\|^2 - \phi^*(\mathbf{y}).$$

The identity $\phi = \phi^{**}$ then gives

$$\frac{1}{2}\|\mathbf{x}\|^2 - \phi(\mathbf{x}) = \inf_{\mathbf{y}} \left\{ \frac{1}{2}\|\mathbf{x} - \mathbf{y}\|^2 - \psi(\mathbf{y}) \right\},$$

which proves the converse. Moreover,

$$\varphi(\mathbf{x}) + \varphi^c(\mathbf{y}) = \frac{1}{2}\|\mathbf{x} - \mathbf{y}\|^2$$

is equivalent to the **Fenchel equality**

$$\phi(\mathbf{x}) + \phi^*(\mathbf{y}) = \mathbf{x} \cdot \mathbf{y},$$

which holds if and only if $\mathbf{y} \in \partial\phi(\mathbf{x})$. □

The Fundamental Optimality Theorem

Theorem 1.7 Let $c: X \times Y \rightarrow \mathbb{R}$ be continuous and bounded below. Suppose that there exist $a \in L^1(\mu)$ and $b \in L^1(\nu)$ such that

$$c(x, y) \leq a(x) + b(y). \tag{1.3}$$

For $\gamma \in \Pi(\mu, \nu)$, the following statements are equivalent:

1. γ is optimal for the Monge–Kantorovich problem;
2. $\text{supp}(\gamma)$ is c -cyclically monotone;
3. there exists a c -concave function φ such that $\max\{\varphi, 0\} \in L^1(\mu)$ and

$$\text{supp}(\gamma) \subseteq \partial^c \varphi.$$

Proof. (1) *implies* (2). Suppose instead that $\text{supp}(\gamma)$ is not c -cyclically monotone. Then there are $(x_i, y_i)_{i=1}^N \subseteq \text{supp}(\gamma)$ and a cyclic permutation σ for which the inequality in (1.2) is strict in the opposite direction. Continuity and the Hausdorff property give neighbourhoods $U_i \times V_i$ on which the same strict improvement persists; they may be chosen disjoint for distinct pairs, while repeated occurrences use the same neighbourhood.

Let γ_i be the normalization of the restriction of γ to $U_i \times V_i$, and let μ_i and ν_i be its marginals.

Choose $\varepsilon > 0$ small enough that subtracting all occurrences of every restricted measure preserves nonnegativity, and define

$$\tilde{\gamma} = \gamma - \varepsilon \sum_{i=1}^N \gamma_i + \varepsilon \sum_{i=1}^N \mu_i \otimes \nu_{\sigma(i)}.$$

The two signed perturbations have the same marginals, so $\tilde{\gamma} \in \Pi(\mu, \nu)$. The strict local cost inequality makes $\tilde{\gamma}$ cheaper than γ , contradicting optimality.

(2) *implies* (3). Set $\Gamma = \text{supp}(\gamma)$ and fix $(\bar{x}, \bar{y}) \in \Gamma$. For $x \in X$, define

$$\begin{aligned} \varphi(x) = \inf \Big\{ & c(x, y_N) - c(x_N, y_N) \\ & + \sum_{i=1}^{N-1} (c(x_{i+1}, y_i) - c(x_i, y_i)) : N \geq 1, (x_i, y_i) \in \Gamma, (x_1, y_1) = (\bar{x}, \bar{y}) \Big\}. \end{aligned} \quad (1.4)$$

Cyclical monotonicity implies $\varphi(\bar{x}) = 0$, so the function is proper. Each expression inside the infimum has the form $c(x, y_N) - A$; hence φ is c -concave. Appending a point $(x, y) \in \Gamma$ to a chain in (1.4) yields

$$\varphi(x') \leq \varphi(x) + c(x', y) - c(x, y)$$

for every $x' \in X$. This is equivalent to $(x, y) \in \partial^c \varphi$, and therefore $\Gamma \subseteq \partial^c \varphi$. The majorant in (1.3), together with an additive normalization of φ , gives $\max\{\varphi, 0\} \in L^1(\mu)$.

(3) *implies* (1). For every coupling $\tilde{\gamma} \in \Pi(\mu, \nu)$,

$$c(x, y) \geq \varphi(x) + \varphi^c(y).$$

Equality holds γ -almost everywhere because $\text{supp}(\gamma) \subseteq \partial^c \varphi$. Consequently,

$$\begin{aligned} \int c \, d\gamma &= \int \varphi \, d\mu + \int \varphi^c \, d\nu \\ &= \int (\varphi(x) + \varphi^c(y)) \, d\tilde{\gamma}(x, y) \leq \int c \, d\tilde{\gamma}. \end{aligned}$$

Thus γ is optimal. □

The theorem shows that optimality depends on the geometry of the support rather than on the particular amount of mass assigned to each supported pair.

Exercise 1.8 Let γ be optimal, and let $A \subseteq X \times Y$ satisfy $\gamma(A) > 0$. Show that the normalized restriction

$$\tilde{\gamma} = \frac{1}{\gamma(A)} \gamma|_A$$

is optimal with respect to its own marginals.

Solution. Let $\tilde{\mu} = (\pi_X)_\# \tilde{\gamma}$ and $\tilde{\nu} = (\pi_Y)_\# \tilde{\gamma}$. For every Borel set $B \subseteq X$,

$$\tilde{\mu}(B) = \frac{\gamma(A \cap (B \times Y))}{\gamma(A)} \leq \frac{\mu(B)}{\gamma(A)},$$

and similarly $\tilde{\nu}(C) \leq \nu(C)/\gamma(A)$ for every Borel set $C \subseteq Y$. Consequently, the functions in the integrable majorant (1.3) satisfy

$$\int_X |a| d\tilde{\mu} \leq \frac{1}{\gamma(A)} \int_X |a| d\mu < \infty \quad \text{and} \quad \int_Y |b| d\tilde{\nu} \leq \frac{1}{\gamma(A)} \int_Y |b| d\nu < \infty.$$

The implication from (1) to (2) in Theorem 1.7 shows that $\text{supp}(\gamma)$ is c -cyclically monotone. Since

$$\text{supp}(\tilde{\gamma}) \subseteq \text{supp}(\gamma),$$

the support of $\tilde{\gamma}$ is also c -cyclically monotone. The implication from (2) to (1) in Theorem 1.7, applied to the marginals $\tilde{\mu}$ and $\tilde{\nu}$, now shows that $\tilde{\gamma}$ is optimal. $\square$

Exercise 1.9 Let $T: X \rightarrow Y$ be measurable and have a c -cyclically monotone graph. Given $\mu \in \mathcal{P}(X)$, set $\nu = T_\# \mu$. Show, under the integrability condition of Theorem 1.7, that T is an optimal transport map from μ to ν .

Solution. Define

$$\gamma = (\text{Id}, T)_\# \mu.$$

Its marginals are μ and $T_\# \mu = \nu$, so $\gamma \in \Pi(\mu, \nu)$. Let G be the graph of T . Since the map $x \mapsto (x, T(x))$ takes values in G , $\gamma(X \times Y \setminus \overline{G}) = 0$, and hence

$$\text{supp}(\gamma) \subseteq \overline{G}.$$

Moreover, $\overline{G}$ remains c -cyclically monotone. Indeed, for any finite collection $(x_i, y_i)_{i=1}^N \subseteq \overline{G}$, choose $(x_i^{(k)}, y_i^{(k)}) \in G$ converging to (x_i, y_i) for each i . The cyclical monotonicity inequality holds for the

points $(x_i^{(k)}, y_i^{(k)})$, and continuity of c allows us to pass to the limit. Thus $\text{supp}(\gamma)$ is c -cyclically monotone.

Under the integrability hypothesis, the implication from (2) to (1) in [Theorem 1.7](#) shows that γ is an optimal coupling. Since γ is induced by T , this says precisely that T is an optimal transport map from μ to ν . $\square$

Kantorovich Duality

Theorem 1.10 (Kantorovich duality) Under the hypotheses of [Theorem 1.7](#),

$$\inf_{\gamma \in \Pi(\mu, \nu)} \int c \, d\gamma = \sup_{\substack{\varphi \in L^1(\mu), \psi \in L^1(\nu) \\ \varphi(x) + \psi(y) \leq c(x, y)}} \left\{ \int \varphi \, d\mu + \int \psi \, d\nu \right\}. \quad (1.5)$$

The supremum is attained by a pair (φ, φ^c) in which φ is c -concave. Such a pair is called a pair of **Kantorovich potentials**.

Proof. For every feasible dual pair and every coupling γ ,

$$\int \varphi \, d\mu + \int \psi \, d\nu = \int (\varphi(x) + \psi(y)) \, d\gamma(x, y) \leq \int c \, d\gamma.$$

This is **weak duality**. By [Theorem 1.3](#), there is an optimal coupling γ^* . The implication from (1) to (3) in [Theorem 1.7](#) gives a c -concave φ with $\text{supp}(\gamma^*) \subseteq \partial^c \varphi$. Equality therefore holds in the pointwise dual constraint γ^* -almost everywhere, and

$$\int c \, d\gamma^* = \int \varphi \, d\mu + \int \varphi^c \, d\nu.$$

This proves equality of the primal and dual values and attainment of the dual value. $\square$

The dual variables have the economic interpretation of prices at the source and destination. **Complementary slackness** is the equality

$$\varphi(x) + \varphi^c(y) = c(x, y)$$

on the support of an optimal coupling.

Lecture 2 — Friday, September 20, 2019

Quadratic Transport and Brenier's Theorem

The quadratic cost

$$c(\mathbf{x}, \mathbf{y}) = \frac{1}{2} \|\mathbf{x} - \mathbf{y}\|^2, \quad X = Y = \mathbb{R}^n,$$

turns the convex analytic correspondence in [Proposition 1.6](#) into an explicit transport map.

Theorem 2.1 (Brenier's theorem) Let $\mu, \nu \in \mathcal{P}_2(\mathbb{R}^n)$, and suppose that μ is absolutely continuous with respect to Lebesgue measure. There is a lower semicontinuous convex function $\phi: \mathbb{R}^n \rightarrow \mathbb{R} \cup \{\infty\}$ such that

$$T(\mathbf{x}) = \nabla \phi(\mathbf{x})$$

is an optimal transport map from μ to ν . The optimal transport map is unique μ -almost everywhere.

Proof. We prove existence. Let γ be an optimal coupling. By [Theorem 1.7](#), there is a c -concave function φ such that

$$\Gamma := \text{supp}(\gamma) \subseteq \partial^c \varphi.$$

By [Proposition 1.6](#),

$$\phi(\mathbf{x}) = \frac{1}{2} \|\mathbf{x}\|^2 - \varphi(\mathbf{x})$$

is lower semicontinuous and convex, and $\Gamma \subseteq \partial \phi$. The fibre $\partial \phi(\mathbf{x})$ is nonempty for μ -almost every $\mathbf{x}$, so these points lie in $\text{Dom}(\phi)$. A lower semicontinuous convex function is locally Lipschitz, and hence differentiable almost everywhere, in the interior of its effective domain. Rademacher's theorem and the absolute continuity of μ therefore imply that $\nabla \phi$ is defined μ -almost everywhere. At every differentiability point,

$$\partial \phi(\mathbf{x}) = \{\nabla \phi(\mathbf{x})\}.$$

Thus γ is concentrated on the graph of $\nabla \phi$, proving that $T = \nabla \phi$ is an optimal map.

It remains to justify the two details used above. The boundary of the effective domain of a convex function has Lebesgue measure zero unless that domain is contained in a lower dimensional affine subspace; the latter alternative is impossible here because $\partial \phi(\mathbf{x})$ is nonempty for μ -almost every

$\mathbf{x}$ and μ is absolutely continuous. Thus μ assigns full mass to points in the interior of the effective domain at which ϕ is differentiable.

For uniqueness, let γ_1 and γ_2 be optimal couplings. Their average is also optimal. Applying the preceding argument to $(\gamma_1 + \gamma_2)/2$ shows that this average is concentrated on the graph of a single gradient. Hence both γ_1 and γ_2 are concentrated on that same graph and must agree. In particular, every optimal coupling is induced by the unique map $T = \nabla\phi$, up to a μ -null set. $\square$

In one dimension, the gradient of a convex function is nondecreasing. Consequently, [Theorem 2.1](#) reduces to the **monotone coupling**.

Theorem 2.2 Let $c(\mathbf{x}, \mathbf{y}) = h(\mathbf{x} - \mathbf{y})$ on $\mathbb{R}^n$, where h is strictly convex and satisfies the standard coercivity and regularity conditions. If μ is absolutely continuous, then there is a unique optimal transport map,

$$T(\mathbf{x}) = \mathbf{x} - \nabla h^*(\nabla\varphi(\mathbf{x})),$$

where φ is a c -concave potential.

The same strategy extends the cost $d^2/2$ to Riemannian manifolds. These extensions are developed in [Gangbo and McCann \(1996\)](#) and [McCann \(2001\)](#).

The regularity step rests on a useful stability property of **moduli of continuity**.

Lemma 2.3 Suppose that every member of a family $(f_\alpha)_\alpha$ satisfies

$$|f_\alpha(x) - f_\alpha(x')| \leq \omega(d(x, x'))$$

for the same modulus of continuity ω . Then $f = \inf_\alpha f_\alpha$ satisfies the same estimate.

Proof. For every $\varepsilon > 0$, choose α with $f_\alpha(x') \leq f(x') + \varepsilon$. Then

$$f(x) - f(x') \leq f_\alpha(x) - f_\alpha(x') + \varepsilon \leq \omega(d(x, x')) + \varepsilon.$$

Let $\varepsilon \downarrow 0$ and interchange x and x' . $\square$

If smooth positive densities f and g and a smooth orientation-preserving diffeomorphism T satisfy

$T_{\#}(f \, d\mathbf{x}) = g \, d\mathbf{y}$, then the change of variables formula gives

$$f(\mathbf{x}) = g(T(\mathbf{x})) |\det DT(\mathbf{x})|.$$

With $T = \nabla\phi$, this becomes the **Monge–Ampère equation**

$$f(\mathbf{x}) = g(\nabla\phi(\mathbf{x})) \det D^2\phi(\mathbf{x}). \quad (2.1)$$

This identity holds at points where the classical change of variables formula applies. Without this smoothness, the equation must be interpreted in a weak sense. Its regularity theory is nonlinear and delicate; for general costs, a central role is played by the **Ma–Trudinger–Wang condition**.

The Wasserstein Metric

Definition 2.4 Let (X, d) be a Polish metric space and let $p \geq 1$. The **p -Wasserstein space** is

$$\mathcal{P}_p(X) = \left\{ \mu \in \mathcal{P}(X) : \int_X d(x_0, x)^p \, d\mu(x) < \infty \text{ for some } x_0 \in X \right\}.$$

For $\mu, \nu \in \mathcal{P}_p(X)$, the **p -Wasserstein distance** is

$$W_p(\mu, \nu) = \left(\inf_{\gamma \in \Pi(\mu, \nu)} \int_{X \times X} d(x, y)^p \, d\gamma(x, y) \right)^{1/p}. \quad (2.2)$$

For point masses, $W_p(\delta_x, \delta_y) = d(x, y)$. Hence $x \mapsto \delta_x$ is an isometric embedding of (X, d) into $(\mathcal{P}_p(X), W_p)$.

Theorem 2.5 The function W_p is a metric on $\mathcal{P}_p(X)$.

Proof. Only the triangle inequality requires a new argument. Let $\mu_1, \mu_2, \mu_3 \in \mathcal{P}_p(X)$, and choose optimal couplings $\gamma_{12} \in \Pi(\mu_1, \mu_2)$ and $\gamma_{23} \in \Pi(\mu_2, \mu_3)$. Disintegrate them over their common marginal μ_2 and define

$$d\gamma_{123}(x_1, x_2, x_3) = d\mu_2(x_2) \, d\gamma_{12}(x_1 \mid x_2) \, d\gamma_{23}(x_3 \mid x_2).$$

This **gluing construction** has $(1, 2)$ -marginal γ_{12} and $(2, 3)$ -marginal γ_{23} . Its $(1, 3)$ -marginal is a feasible coupling of μ_1 and μ_3 . The triangle inequality for d followed by Minkowski's inequality

gives

$$\begin{aligned} W_p(\mu_1, \mu_3) &\leq \left(\int d(x_1, x_3)^p d\gamma_{123} \right)^{1/p} \\ &\leq \left(\int d(x_1, x_2)^p d\gamma_{123} \right)^{1/p} + \left(\int d(x_2, x_3)^p d\gamma_{123} \right)^{1/p} \\ &= W_p(\mu_1, \mu_2) + W_p(\mu_2, \mu_3). \end{aligned}$$

□

For Gaussian measures the quadratic Wasserstein metric separates location and dispersion.

Proposition 2.6 Let $\mu = \mathcal{N}(\mathbf{m}_\alpha, \Sigma_\alpha)$ and $\nu = \mathcal{N}(\mathbf{m}_\beta, \Sigma_\beta)$. Then

$$W_2(\mu, \nu)^2 = \|\mathbf{m}_\alpha - \mathbf{m}_\beta\|^2 + B(\Sigma_\alpha, \Sigma_\beta)^2,$$

where the **Bures–Wasserstein distance** is

$$B(\Sigma_\alpha, \Sigma_\beta)^2 = \text{tr} \left(\Sigma_\alpha + \Sigma_\beta - 2(\Sigma_\alpha^{1/2} \Sigma_\beta \Sigma_\alpha^{1/2})^{1/2} \right).$$

Dirichlet Transport and Divergences

The open simplex

$$\Delta_n = \left\{ \mathbf{p} = (p_1, \dots, p_n) \in (0, 1)^n : \sum_{i=1}^n p_i = 1 \right\}$$

supports a multiplicative analogue of quadratic transport. Its cost is

$$c(\mathbf{p}, \mathbf{q}) = \log \left(\frac{1}{n} \sum_{i=1}^n \frac{q_i}{p_i} \right) - \frac{1}{n} \sum_{i=1}^n \log \left(\frac{q_i}{p_i} \right). \quad (2.3)$$

The expression is nonnegative by the arithmetic–geometric mean inequality. It is a **divergence** rather than a metric: symmetry and the triangle inequality need not hold.

One motivation comes from stochastic portfolio theory. If $m_i(t)$ is the market capitalization of stock i divided by the total market capitalization, then $\mathbf{m}(t) \in \Delta_n$. A positive concave function

Φ measures market diversity. Typical examples are

$$\Phi(\mathbf{p}) = \left(\sum_{i=1}^n p_i^\lambda \right)^{1/\lambda}, \quad 0 < \lambda < 1 \quad \text{and} \quad \Phi(\mathbf{p}) = - \sum_{i=1}^n p_i \log p_i.$$

The first is a **diversity index**; the second is **Shannon entropy**. Stability of the capital distribution can therefore be phrased as persistence of diversity along the market path.

A positive differentiable concave function Φ generates a portfolio map $\boldsymbol{\pi}: \Delta_n \rightarrow \bar{\Delta}_n$ by

$$\pi_i(\mathbf{p}) = p_i (1 + D_{\mathbf{e}_i - \mathbf{p}} \log \Phi(\mathbf{p})). \quad (2.4)$$

Geometrically, the supporting hyperplane to Φ at $\mathbf{p}$ determines the portfolio weights, as shown in Figure 2.

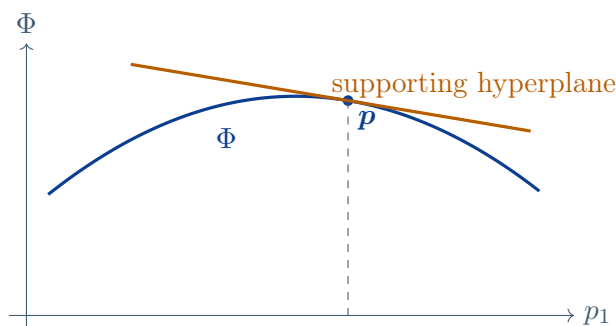


FIGURE 2

Two useful examples follow directly from (2.4). If $\boldsymbol{\pi} \in \bar{\Delta}_n$ and

$$\Phi(\mathbf{p}) = \prod_{i=1}^n p_i^{\pi_i},$$

then the generated portfolio is the constant weight vector $\boldsymbol{\pi}$. If

$$\Phi(\mathbf{p}) = \left(\sum_{i=1}^n p_i^\lambda \right)^{1/\lambda},$$

then $\pi_i(\mathbf{p})$ is proportional to p_i^λ .

Theorem 2.7 For a portfolio map $\boldsymbol{\pi}: \Delta_n \rightarrow \bar{\Delta}_n$, the following statements are equivalent:

1. The portfolio is functionally generated.
2. For every cycle $\mathbf{p}(0), \mathbf{p}(1), \dots, \mathbf{p}(T) = \mathbf{p}(0)$,

$$\prod_{t=0}^{T-1} \left(\sum_{i=1}^n \pi_i(\mathbf{p}(t)) \frac{p_i(t+1)}{p_i(t)} \right) \geq 1.$$

Condition (2) says that the portfolio does not underperform the market over a closed path. After a change of coordinates, it is the c -cyclical monotonicity condition for the cost in (2.3).

The relevant multiplicative linear structure on Δ_n is defined by

$$(\mathbf{p} \oplus \mathbf{q})_i = \frac{p_i q_i}{\sum_{j=1}^n p_j q_j} \quad \text{and} \quad (\lambda \odot \mathbf{p})_i = \frac{p_i^\lambda}{\sum_{j=1}^n p_j^\lambda}.$$

The neutral element is $\mathbf{e} = (1/n, \dots, 1/n)$, and the inverse $\mathbf{p}^{-1}$ is obtained by normalizing the reciprocals. This is **Aitchison's geometry** for compositional data.

Lemma 2.8 The cost in (2.3) satisfies

$$c(\mathbf{p}, \mathbf{q}) = \text{KL}(\mathbf{e} \mid \mathbf{q} \oplus \mathbf{p}^{-1}),$$

where the **relative entropy** is

$$\text{KL}(\mathbf{r} \mid \mathbf{s}) = \sum_{i=1}^n r_i \log \left(\frac{r_i}{s_i} \right).$$

Thus the **Dirichlet cost** plays the role that squared Euclidean distance plays in quadratic transport.

Theorem 2.9 Suppose that $\mu \in \mathcal{P}(\Delta_n)$ is absolutely continuous and that μ and ν satisfy the required moment conditions. There is a concave function $\Phi: \Delta_n \rightarrow (0, \infty)$ for which the deterministic optimal transport map has components

$$q_i = \frac{p_i \pi_i(\mathbf{p}^{-1})}{\sum_{j=1}^n p_j \pi_j(\mathbf{p}^{-1})}.$$

Equivalently, $\mathbf{q} = \mathbf{p} \oplus \pi(\mathbf{p}^{-1})$.

The increment $\mathbf{q} \oplus \mathbf{p}^{-1}$ is the multiplicative counterpart of the displacement $\mathbf{y} - \mathbf{x}$.

Definition 2.10 Let $\Omega \subseteq \mathbb{R}^n$ be convex, and let $\phi: \Omega \rightarrow \mathbb{R}$ be differentiable and strictly convex. The **Bregman divergence** generated by ϕ is

$$D_\phi[\mathbf{x} : \mathbf{x}'] = \phi(\mathbf{x}) - \phi(\mathbf{x}') - \nabla \phi(\mathbf{x}') \cdot (\mathbf{x} - \mathbf{x}').$$

Convexity gives $D_\phi[\mathbf{x} : \mathbf{x}'] \geq 0$, with equality only when $\mathbf{x} = \mathbf{x}'$. The divergence need not be symmetric. If ϕ is twice differentiable at $\mathbf{x}$, then locally,

$$D_\phi[\mathbf{x} + \Delta \mathbf{x} : \mathbf{x}] = \frac{1}{2} \Delta \mathbf{x}^\top D^2 \phi(\mathbf{x}) \Delta \mathbf{x} + o(\|\Delta \mathbf{x}\|^2),$$

so the Hessian of ϕ determines a Riemannian metric. The quadratic generator $\phi(\mathbf{x}) = \|\mathbf{x}\|^2/2$ gives squared Euclidean distance, while negative Shannon entropy gives relative entropy.

Definition 2.11 Let T be an optimal transport map with graph $G = \{(\mathbf{p}, T(\mathbf{p}))\}$, and let (φ, ψ) be Kantorovich potentials satisfying

$$\varphi(\mathbf{p}) + \psi(\mathbf{q}') \leq c(\mathbf{p}, \mathbf{q}'),$$

with equality precisely when $\mathbf{q}' = T(\mathbf{p})$. For $x = (\mathbf{p}, \mathbf{q})$ and $x' = (\mathbf{p}', \mathbf{q}')$ in G , define the **c-divergence**

$$D[x : x'] = c(\mathbf{p}, \mathbf{q}') - \varphi(\mathbf{p}) - \psi(\mathbf{q}') \geq 0.$$

The crossed pair $(\mathbf{p}, \mathbf{q}')$ generally lies off the optimal graph, as illustrated in Figure 3; the c -divergence measures its **dual slack**.

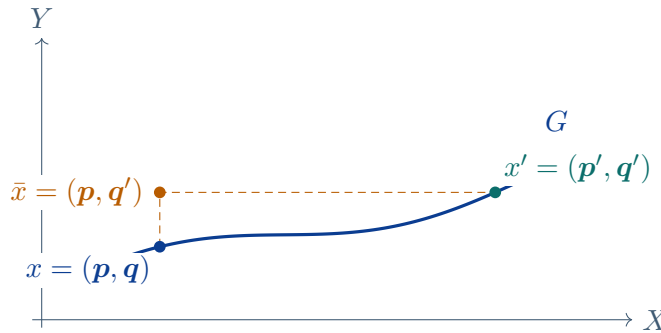


FIGURE 3

Theorem 2.12 For $c(\mathbf{x}, \mathbf{y}) = \|\mathbf{x} - \mathbf{y}\|^2/2$, the c -divergence is a Bregman divergence.

Proof. Apply the convex correspondence in [Proposition 1.6](#). The equality set $\partial^c \varphi$ becomes the ordinary subdifferential of the convex function $\phi = \|\cdot\|^2/2 - \varphi$, and the dual slack becomes

$$\phi(\mathbf{x}) - \phi(\mathbf{x}') - \nabla \phi(\mathbf{x}') \cdot (\mathbf{x} - \mathbf{x}').$$

□

Definition 2.13 Let $\alpha > 0$. A differentiable function $\varphi: \Omega \rightarrow \mathbb{R}$ is **α -exponentially concave** if $\Phi = e^{\alpha\varphi}$ is concave. Its **$L^{(\alpha)}$ -divergence** is

$$D^{(\alpha)}[\xi : \xi'] = \frac{1}{\alpha} \log (1 + \alpha \nabla \varphi(\xi') \cdot (\xi - \xi')) - (\varphi(\xi) - \varphi(\xi')).$$

The first term uses a logarithmic approximation to the graph of φ , while a Bregman divergence uses a linear approximation. As $\alpha \downarrow 0$, the logarithmic approximation becomes linear and $D^{(\alpha)}$ converges to the corresponding Bregman divergence. The Dirichlet transport is the case $\Omega = \Delta_n$ and $\alpha = 1$.

Wasserstein Convergence and Displacement Interpolation

Theorem 2.14 Let $(\mu_k)_{k \geq 1} \subseteq \mathcal{P}_p(X)$ and $\mu \in \mathcal{P}_p(X)$. The following statements are equivalent:

1. $W_p(\mu_k, \mu) \rightarrow 0$;
2. $\mu_k \Rightarrow \mu$ and, for some, hence every, $x_0 \in X$,

$$\int_X d(x_0, x)^p d\mu_k(x) \longrightarrow \int_X d(x_0, x)^p d\mu(x).$$

Proof. We prove the result when X is compact; the general case replaces boundedness by uniform integrability.

Suppose first that (2) holds. The moment condition is automatic on a compact space. By the

Skorokhod representation theorem, there are random elements $X_k \sim \mu_k$ and $X \sim \mu$ on a common probability space with $X_k \rightarrow X$ almost surely. Their joint laws are couplings, and the bounded convergence theorem gives

$$W_p(\mu_k, \mu)^p \leq \mathbb{E}[d(X_k, X)^p] \rightarrow 0.$$

Conversely, suppose that (1) holds and let (X_k, X) be an optimal coupling of (μ_k, μ) . For $\delta > 0$, Markov's inequality gives

$$\mathbb{P}(d(X_k, X) \geq \delta) \leq \frac{W_p(\mu_k, \mu)^p}{\delta^p}.$$

If f is continuous, compactness makes it bounded by some M and uniformly continuous. Given $\varepsilon > 0$, choose $\delta > 0$ such that $d(x, x') < \delta$ implies $|f(x) - f(x')| < \varepsilon$. Then

$$\begin{aligned} \left| \int f d\mu_k - \int f d\mu \right| &\leq \mathbb{E}[|f(X_k) - f(X)|] \\ &\leq \varepsilon + 2M \frac{W_p(\mu_k, \mu)^p}{\delta^p}. \end{aligned}$$

Letting $k \rightarrow \infty$ and then $\varepsilon \downarrow 0$ proves weak convergence. The moment convergence again follows automatically from compactness. $\square$

Definition 2.15 Let $\mu_0, \mu_1 \in \mathcal{P}_2(\mathbb{R}^n)$, with μ_0 absolutely continuous, and let $\nabla\phi$ be the **Brenier map** satisfying $\mu_1 = (\nabla\phi)_\# \mu_0$. Define

$$\phi_t(\mathbf{x}) = (1-t)\frac{1}{2}\|\mathbf{x}\|^2 + t\phi(\mathbf{x})$$

and the **displacement interpolation**

$$[\mu_0, \mu_1]_t = (\nabla\phi_t)_\# \mu_0 = ((1-t)\text{Id} + t\nabla\phi)_\# \mu_0. \quad (2.5)$$

Ordinary mixture interpolation fades one density out while fading the other in. Displacement interpolation instead moves individual particles along straight lines. [Figure 4](#) shows how two separated masses travel rather than coexist at intermediate times.

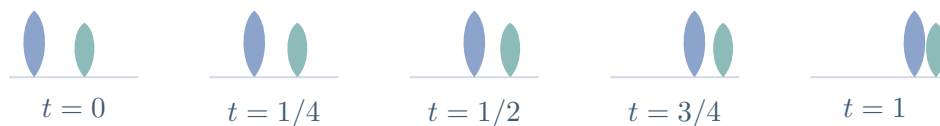


FIGURE 4

Proposition 2.16 Let $\mu, \nu \in \mathcal{P}_{2,\text{ac}}(\mathbb{R}^n)$. Then the following hold:

1. $[\mu, \nu]_t \in \mathcal{P}_{2,\text{ac}}(\mathbb{R}^n)$ for $0 < t < 1$;
2. $[\mu, \nu]_t = [\nu, \mu]_{1-t}$;
3. $[[\mu, \nu]_t, [\mu, \nu]_{t'}]_s = [\mu, \nu]_{(1-s)t+st'}$;
4. For μ -almost every $\mathbf{x}$, the path $t \mapsto (1-t)\mathbf{x} + t\nabla\phi(\mathbf{x})$ is a constant velocity straight line.

Theorem 2.17 The displacement interpolation is a **constant speed geodesic**:

$$W_2([\mu, \nu]_s, [\mu, \nu]_t) = |s - t|W_2(\mu, \nu).$$

Proof. Write $\mu_t = [\mu_0, \mu_1]_t$. Coupling μ_s and μ_t through the same starting point gives

$$\begin{aligned} W_2(\mu_s, \mu_t)^2 &\leq \int \|\nabla\phi_s(\mathbf{x}) - \nabla\phi_t(\mathbf{x})\|^2 d\mu_0(\mathbf{x}) \\ &= (s - t)^2 \int \|\nabla\phi(\mathbf{x}) - \mathbf{x}\|^2 d\mu_0(\mathbf{x}) \\ &= (s - t)^2 W_2(\mu_0, \mu_1)^2. \end{aligned}$$

Applying this estimate to the segments from 0 to s , from s to t , and from t to 1, the triangle inequality forces equality. $\square$

Definition 2.18 A set $A \subseteq \mathcal{P}_2(\mathbb{R}^n)$ is **displacement convex** if it contains $[\mu, \nu]_t$ whenever it contains μ and ν and the Brenier interpolation in (2.5) is defined. A functional $F: A \rightarrow \mathbb{R}$ is **displacement convex** if

$$t \mapsto F([\mu, \nu]_t)$$

is convex for every such pair $\mu, \nu \in A$.

This convexity was introduced to study interacting gases. If ρ is a density, consider

$$E(\rho) = \int A(\rho(\mathbf{x})) d\mathbf{x} + \frac{1}{2} \iint V(\mathbf{x} - \mathbf{y}) \rho(\mathbf{x}) \rho(\mathbf{y}) d\mathbf{x} d\mathbf{y}.$$

The terms represent **internal compression energy** and **interaction energy**. Although E need not be convex under ordinary mixtures of densities, it is displacement convex under natural

physical conditions on A and V .

The same geodesic picture survives beyond Euclidean space.

Theorem 2.19 If (X, d) is Polish and geodesic, then $(\mathcal{P}_2(X), W_2)$ is geodesic. A curve $t \mapsto \mu_t$ is a constant speed geodesic if and only if there is a measure $\mathbb{P} \in \mathcal{P}_2(\text{Geod}(X))$ whose marginal at time t is μ_t and whose marginal at every pair of times is an optimal coupling.

The measure $\mathbb{P}$ describes a random particle whose path is itself a geodesic in X .

Dynamic Optimal Transport

In the static problem, a pair (x, y) records only the endpoints. A **dynamic cost** assigns an **action** $C[(z_t)_{0 \leq t \leq 1}]$ to a path, with compatibility condition

$$c(x, y) = \inf \{C[(z_t)_{0 \leq t \leq 1}] : z_0 = x, z_1 = y\}.$$

For the quadratic cost, the basic identity is

$$\frac{1}{2} \|x - y\|^2 = \inf_{\substack{x_0 = x \\ x_1 = y \\ x. \text{ absolutely continuous}}} \int_0^1 \frac{1}{2} \|\dot{x}_t\|^2 dt. \quad (2.6)$$

The unique minimizing path is the constant velocity line segment.

In the **Lagrangian description**, one follows individual particles. If $\mathcal{D}_1(\mu, \nu)$ is the set of laws on path space with endpoint marginals μ and ν , then

$$\frac{1}{2} W_2(\mu, \nu)^2 = \inf_{\mathbb{P} \in \mathcal{D}_1(\mu, \nu)} \mathbb{E}_{\mathbb{P}} \left[\int_0^1 \frac{1}{2} \|\dot{X}_t\|^2 dt \right]. \quad (2.7)$$

An optimizer is obtained by first drawing an optimally coupled endpoint pair and then joining it by the line in (2.6).

The **Eulerian description** instead follows the collective density. Let $v(t, x)$ be a **velocity field** and let μ_t be the distribution of particles solving

$$\dot{x}_t = v(t, x_t).$$

Lemma 2.20 Let $(\mathbf{X}_t)_{0 \leq t \leq 1}$ be an absolutely continuous random path satisfying

$$\dot{\mathbf{X}}_t = \mathbf{v}(t, \mathbf{X}_t)$$

for almost every t , and assume that $\mathbb{E} \int_0^1 \|\mathbf{v}(t, \mathbf{X}_t)\| dt < \infty$. Then its marginal laws $\mu_t = \mathcal{L}(\mathbf{X}_t)$ satisfy the **continuity equation**

$$\partial_t \mu_t + \nabla \cdot (\mu_t \mathbf{v}) = 0 \quad (2.8)$$

in the weak sense.

Proof. For every smooth ζ with compact spatial support on $[0, 1] \times \mathbb{R}^n$, the chain rule along almost every path gives

$$\zeta(1, \mathbf{X}_1) - \zeta(0, \mathbf{X}_0) = \int_0^1 \{ \partial_t \zeta(t, \mathbf{X}_t) + \nabla \zeta(t, \mathbf{X}_t) \cdot \mathbf{v}(t, \mathbf{X}_t) \} dt.$$

The integrability hypothesis permits expectation and time integration to be interchanged. Rewriting the expectations using the marginal laws yields

$$\int_0^1 \int_{\mathbb{R}^n} (\partial_t \zeta + \nabla \zeta \cdot \mathbf{v}) d\mu_t dt = \int_{\mathbb{R}^n} \zeta(1, \mathbf{x}) d\mu_1(\mathbf{x}) - \int_{\mathbb{R}^n} \zeta(0, \mathbf{x}) d\mu_0(\mathbf{x}),$$

which is the weak form of (2.8). □

Theorem 2.21 (Benamou–Brenier formula) For $\mu, \nu \in \mathcal{P}_2(\mathbb{R}^n)$,

$$\frac{1}{2} W_2(\mu, \nu)^2 = \inf_{(\rho, \mathbf{v})} \int_0^1 \int_{\mathbb{R}^n} \frac{1}{2} \|\mathbf{v}(t, \mathbf{x})\|^2 d\rho_t(\mathbf{x}) dt, \quad (2.9)$$

where the infimum is over solutions of (2.8) with $\rho_0 = \mu$ and $\rho_1 = \nu$.

Here the admissible curves are narrowly continuous, the velocity field is measurable, and the displayed action is finite; the continuity equation is interpreted distributionally.

The Lagrangian and Eulerian pictures are two views of the same motion: one resolves particle trajectories, while the other records only the density and velocity field.

Empirical Wasserstein Distances

Given independent samples $X_1, \dots, X_N \sim \mu$ and $Y_1, \dots, Y_N \sim \nu$, define the **empirical measures**

$$\mu_N = \frac{1}{N} \sum_{i=1}^N \delta_{X_i} \quad \text{and} \quad \nu_N = \frac{1}{N} \sum_{i=1}^N \delta_{Y_i}.$$

On bounded d -dimensional domains in the high dimensional regime $d > 4$, a characteristic estimate is

$$\mathbb{E} [|W_2(\mu_N, \nu_N) - W_2(\mu, \nu)|] = O(N^{-1/d})$$

reveals the **curse of dimensionality**. The critical dimension $d = 4$ introduces a logarithmic correction, while below it the dimension-free sampling term dominates the generic moment bound; additional regularity can improve these rates.

More refined bounds use the moment

$$M_q(\mu) = \int_{\mathbb{R}^d} \|\mathbf{x}\|^q d\mu(\mathbf{x}).$$

Theorem 2.22 Let $\mu \in \mathcal{P}(\mathbb{R}^d)$, let $p \geq 1$, and suppose that $M_q(\mu) < \infty$ for some $q > p$. There is a constant C , depending only on p , d , and q , such that

$$\mathbb{E}[W_p(\mu_N, \mu)^p] \leq CM_q(\mu)^{p/q} R_{p,d,q}(N),$$

where

$$R_{p,d,q}(N) = \begin{cases} N^{-1/2} + N^{-(q-p)/q}, & p > d/2 \text{ and } q \neq 2p, \\ N^{-1/2} \log(1+N) + N^{-(q-p)/q}, & p = d/2 \text{ and } q \neq 2p, \\ N^{-p/d} + N^{-(q-p)/q}, & 1 \leq p < d/2 \text{ and } q \neq dp/(d-p). \end{cases}$$

This result is due to [Fournier and Guillin \(2015\)](#). It separates the **geometric rate**, controlled by the dimension, from the **tail rate**, controlled by q .

Lecture 3 — Friday, September 27, 2019

2. Discrete Algorithms and Entropic Regularization

Discrete Transport as a Linear Program

Suppose that μ and ν are supported on finite sets and are represented by probability vectors

$$\mathbf{a} = (a_i)_{i=1}^n \quad \text{and} \quad \mathbf{b} = (b_j)_{j=1}^m.$$

The cost is represented by a matrix $\mathbf{C} = (C_{ij}) \in \mathbb{R}^{n \times m}$. A coupling is now a nonnegative matrix $\mathbf{P} = (P_{ij})$ in the **transport polytope**

$$\mathcal{U}(\mathbf{a}, \mathbf{b}) = \left\{ \mathbf{P} \in \mathbb{R}_+^{n \times m} : \mathbf{P} \mathbf{1}_m = \mathbf{a}, \mathbf{P}^\top \mathbf{1}_n = \mathbf{b} \right\}. \quad (3.1)$$

Thus the Monge–Kantorovich problem becomes the linear program

$$L_{\mathbf{C}}(\mathbf{a}, \mathbf{b}) = \min_{\mathbf{P} \in \mathcal{U}(\mathbf{a}, \mathbf{b})} \langle \mathbf{P}, \mathbf{C} \rangle = \min_{\mathbf{P} \in \mathcal{U}(\mathbf{a}, \mathbf{b})} \sum_{i=1}^n \sum_{j=1}^m C_{ij} P_{ij}. \quad (3.2)$$

In standard linear programming notation, a primal problem

$$\max_{\mathbf{x}} \mathbf{c}^\top \mathbf{x} \quad \text{subject to} \quad \mathbf{A} \mathbf{x} = \mathbf{b}, \mathbf{x} \geq \mathbf{0}$$

has the dual problem

$$\min_{\mathbf{y}} \mathbf{b}^\top \mathbf{y} \quad \text{subject to} \quad \mathbf{A}^\top \mathbf{y} \geq \mathbf{c}.$$

Applied to (3.2), these constraints reproduce the Kantorovich dual in [Theorem 1.10](#).

The variables also form the bipartite network in [Figure 5](#). A positive entry P_{ij} is a flow from source vertex i to target vertex j' , and its width records the transported mass.

Proposition 3.1 If $\mathbf{P}$ is an **extreme point** of $\mathcal{U}(\mathbf{a}, \mathbf{b})$, then $\mathbf{P}$ has at most $n + m - 1$ nonzero entries.

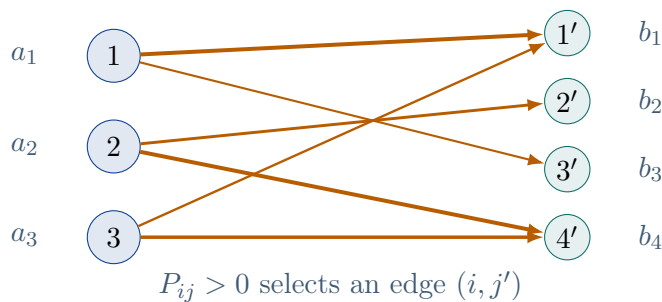


FIGURE 5

Proof. Consider the **support subgraph**

$$\tilde{E} = \{(i, j') : P_{ij} > 0\}.$$

We claim that this graph has no cycle. Suppose that it contains the alternating cycle

$$(i_1, j'_1), (i_2, j'_1), (i_2, j'_2), \dots, (i_k, j'_k), (i_1, j'_k).$$

Construct a matrix $\mathbf{E}$ whose entries alternate between ε and $-\varepsilon$ around this cycle and vanish elsewhere, where

$$0 < \varepsilon < \min\{P_{ij} : (i, j') \in \tilde{E}\}.$$

Every row sum and column sum of $\mathbf{E}$ is zero, so both $\mathbf{P} + \mathbf{E}$ and $\mathbf{P} - \mathbf{E}$ belong to $\mathcal{U}(\mathbf{a}, \mathbf{b})$. Since

$$\mathbf{P} = \frac{1}{2}(\mathbf{P} + \mathbf{E}) + \frac{1}{2}(\mathbf{P} - \mathbf{E}),$$

this contradicts extremality. A forest on $n + m$ vertices has at most $n + m - 1$ edges, proving the result. $\square$

In particular, when $n = m$, an extreme optimal coupling has order n nonzero entries rather than order n^2 . For convex costs on a one dimensional ordered state space, this sparsity appears along the monotone diagonal. Concave costs may instead favour keeping some mass fixed and moving the remainder farther.

Classical solvers exploit the network structure. Scaling versions of the **auction algorithm** have polynomial pseudopolynomial bounds involving the range of the costs. The **network simplex algorithm** is often fast in practice, although common pivot rules can take exponentially many iterations in the worst case; specially designed scaling pivot rules have polynomial guarantees. **Interior point methods** solve more general linear programs but require repeated **Newton systems**. These costs motivate regularization methods whose basic operation is matrix–vector

multiplication.

Entropic Regularization and Sinkhorn Scaling

Throughout this subtopic, assume that every component of $\mathbf{a}$ and $\mathbf{b}$ is strictly positive. Zero-mass rows and columns can be removed before applying the scaling formulas.

Definition 3.2 The **Shannon entropy** of a coupling matrix is

$$H(\mathbf{P}) = - \sum_{i=1}^n \sum_{j=1}^m P_{ij} \log P_{ij}.$$

For $\varepsilon > 0$, the **entropically regularized transport value** is

$$L_{\mathbf{C}}^{\varepsilon}(\mathbf{a}, \mathbf{b}) = \min_{\mathbf{P} \in \mathcal{U}(\mathbf{a}, \mathbf{b})} \{ \langle \mathbf{P}, \mathbf{C} \rangle - \varepsilon H(\mathbf{P}) \}. \quad (3.3)$$

Since $x \mapsto x \log x$, with $0 \log 0 = 0$, is strictly convex on $[0, \infty)$, (3.3) has a unique minimizer. The method was introduced computationally by [Cuturi \(2013\)](#).

For positive vectors of equal total mass, define

$$\text{KL}(\mathbf{p} \mid \mathbf{q}) = \sum_i p_i \log \left(\frac{p_i}{q_i} \right).$$

If $\mathbf{q}$ is uniform, then $\text{KL}(\mathbf{p} \mid \mathbf{q})$ differs from $-H(\mathbf{p})$ by a constant. Thus regularization can be viewed as constraining the relative entropy from the uniform measure. For positive measures of possibly different masses, the extended definition is

$$\text{KL}(\mathbf{a} \mid \mathbf{b}) = \sum_i \left(a_i \log \left(\frac{a_i}{b_i} \right) - a_i + b_i \right). \quad (3.4)$$

With the convention

$$D_f[\mathbf{a} : \mathbf{b}] = \sum_i a_i f \left(\frac{b_i}{a_i} \right),$$

this is the **f -divergence** generated by $f(u) = -\log u + u - 1$.

Let the **Gibbs kernel** be

$$K_{ij} = \exp(-C_{ij}/\varepsilon). \quad (3.5)$$

A direct calculation shows that the objective in (3.3) differs from $\varepsilon \text{KL}(\mathbf{P} \mid \mathbf{K})$ only by constants

that do not depend on $\mathbf{P}$. Therefore

$$\mathbf{P}^\varepsilon = \operatorname{argmin}_{\mathbf{P} \in \mathcal{U}(\mathbf{a}, \mathbf{b})} \operatorname{KL}(\mathbf{P} \mid \mathbf{K}). \quad (3.6)$$

The solution is the **relative entropy projection** of the Gibbs kernel onto the transport polytope.

Theorem 3.3 Let $\mathbf{P}^\varepsilon$ be the unique minimizer in (3.3). Then

$$\lim_{\varepsilon \downarrow 0} \mathbf{P}^\varepsilon = \operatorname{argmax} \{H(\mathbf{P}) : \mathbf{P} \in \mathcal{U}(\mathbf{a}, \mathbf{b}), \langle \mathbf{P}, \mathbf{C} \rangle = L_{\mathbf{C}}(\mathbf{a}, \mathbf{b})\},$$

whereas

$$\lim_{\varepsilon \rightarrow \infty} \mathbf{P}^\varepsilon = \mathbf{a}\mathbf{b}^\top.$$

Proof. Choose a sequence $\varepsilon_\ell \downarrow 0$. Compactness of the transport polytope gives a subsequence for which $\mathbf{P}^{\varepsilon_\ell} \rightarrow \mathbf{P}^*$. If $\mathbf{P}$ is any unregularized optimizer, then optimality of $\mathbf{P}^{\varepsilon_\ell}$ gives

$$0 \leq \langle \mathbf{C}, \mathbf{P}^{\varepsilon_\ell} \rangle - \langle \mathbf{C}, \mathbf{P} \rangle \leq \varepsilon_\ell (H(\mathbf{P}^{\varepsilon_\ell}) - H(\mathbf{P})).$$

Continuity of entropy on the finite transport polytope shows that $\mathbf{P}^*$ is optimal. Dividing the same inequality by ε_ℓ and passing to the limit yields $H(\mathbf{P}) \leq H(\mathbf{P}^*)$. Strict concavity makes the **maximum entropy optimizer** unique, so the entire family converges.

For the second limit, if (X, Y) has joint law $\mathbf{P}$, then

$$I(X; Y) = \operatorname{KL}(\mathbf{P} \mid \mathbf{a}\mathbf{b}^\top) = \sum_{i,j} P_{ij} \log \left(\frac{P_{ij}}{a_i b_j} \right).$$

This **mutual information** is nonnegative and vanishes precisely when X and Y are independent. As the entropy term dominates, the optimizer therefore converges to $\mathbf{a}\mathbf{b}^\top$. $\square$

Theorem 3.4 The minimizer of (3.3) has the form

$$P_{ij}^\varepsilon = u_i v_j K_{ij} \quad \text{and} \quad \mathbf{P}^\varepsilon = \operatorname{diag}(\mathbf{u}) \mathbf{K} \operatorname{diag}(\mathbf{v}), \quad (3.7)$$

for positive vectors $\mathbf{u}$ and $\mathbf{v}$.

Proof. Introduce Lagrange multipliers f_i and g_j for the two marginal constraints. Differentiating

the Lagrangian with respect to P_{ij} gives

$$C_{ij} + \varepsilon(1 + \log P_{ij}) - f_i - g_j = 0.$$

The constant term can be absorbed into either multiplier. Hence

$$P_{ij} = \exp(f_i/\varepsilon) \exp(-C_{ij}/\varepsilon) \exp(g_j/\varepsilon),$$

which has the form (3.7). □

The **scaling vectors** are not individually unique: multiplying $\mathbf{u}$ by a positive constant and dividing $\mathbf{v}$ by the same constant leaves $\mathbf{P}$ unchanged. The marginal constraints become

$$\mathbf{u} \odot (\mathbf{K}\mathbf{v}) = \mathbf{a} \quad \text{and} \quad \mathbf{v} \odot (\mathbf{K}^\top \mathbf{u}) = \mathbf{b},$$

where $\odot$ denotes **componentwise multiplication**.

Algorithm 3.5 (Sinkhorn's algorithm) Choose a positive initial vector $\mathbf{v}^{(0)}$, such as $\mathbf{1}_m$. For $\ell \geq 0$, update componentwise by

$$\mathbf{u}^{(\ell+1)} = \frac{\mathbf{a}}{\mathbf{K}\mathbf{v}^{(\ell)}} \quad \text{and} \quad \mathbf{v}^{(\ell+1)} = \frac{\mathbf{b}}{\mathbf{K}^\top \mathbf{u}^{(\ell+1)}}.$$

Stop when both marginal residuals are within the required tolerance.

Each half step enforces one marginal exactly and perturbs the other. The iterates alternately rescale rows and columns until both constraints agree. The work is dominated by multiplication by $\mathbf{K}$ and $\mathbf{K}^\top$, which is well suited to parallel hardware and to simultaneous processing of several histograms.

Smaller ε produces a sharper, sparser coupling, while larger ε diffuses mass around the low cost region. Other regularizers can also promote sparsity, but entropy has a special probabilistic connection with the Schrödinger bridge.

For a target additive error τ , choosing $\varepsilon = \tau/(4 \log n)$ and applying a final rounding step yields a feasible coupling $\hat{\mathbf{P}}$ with

$$\langle \hat{\mathbf{P}}, \mathbf{C} \rangle \leq L_{\mathbf{C}}(\mathbf{a}, \mathbf{b}) + \tau.$$

A representative bound is

$$O(\|\mathbf{C}\|_\infty^3 \log(n) \tau^{-3})$$

Sinkhorn iterations and $O(n^2 \log(n) \tau^{-3})$ arithmetic operations when the cost normalization is

fixed.

Global convergence can be understood through the **Hilbert projective metric** on the **positive cone**:

$$d_H(\mathbf{u}, \mathbf{u}') = \|\log \mathbf{u} - \log \mathbf{u}'\|_{\text{var}} = \log \max_{i,j} \frac{u_i u'_j}{u_j u'_i}, \quad (3.8)$$

where $\|\mathbf{z}\|_{\text{var}} = \max_i z_i - \min_i z_i$. This is a metric after identifying positive multiples.

Theorem 3.6 (Birkhoff–Samelson contraction) If $\mathbf{K}$ has strictly positive entries, then

$$d_H(\mathbf{K}\mathbf{v}, \mathbf{K}\mathbf{v}') \leq \lambda(\mathbf{K}) d_H(\mathbf{v}, \mathbf{v}'),$$

where $0 \leq \lambda(\mathbf{K}) < 1$ is explicit.

Theorem 3.7 The scaling vectors in [Algorithm 3.5](#) converge in the Hilbert projective metric to scaling vectors $(\mathbf{u}^*, \mathbf{v}^*)$ for the unique regularized coupling.

Proof. Entrywise reciprocal scaling is an isometry in d_H . Combining this fact with [Theorem 3.6](#) gives a strict contraction over a full row–column update. For example,

$$d_H(\mathbf{u}^{(\ell+1)}, \mathbf{u}^*) = d_H(\mathbf{K}\mathbf{v}^{(\ell)}, \mathbf{K}\mathbf{v}^*) \leq \lambda(\mathbf{K}) d_H(\mathbf{v}^{(\ell)}, \mathbf{v}^*).$$

The analogous estimate for $\mathbf{v}$ completes the contraction argument. □

Sinkhorn Divergences and Deconvolution

Even when $C_{ii} = 0$ and $C_{ij} > 0$ for $i \neq j$, the regularized value need not vanish on the diagonal. A debiased version is

$$\tilde{L}_{\mathbf{C}}^{\varepsilon}(\mathbf{a}, \mathbf{b}) = L_{\mathbf{C}}^{\varepsilon}(\mathbf{a}, \mathbf{b}) - \frac{1}{2} L_{\mathbf{C}}^{\varepsilon}(\mathbf{a}, \mathbf{a}) - \frac{1}{2} L_{\mathbf{C}}^{\varepsilon}(\mathbf{b}, \mathbf{b}). \quad (3.9)$$

Under the symmetry, regularity, and positive-kernel hypotheses of the cited result, this **Sinkhorn divergence** is nonnegative and vanishes only when its arguments agree. As the regularization changes, it interpolates between transport geometry and kernel discrepancy geometry; see [Feydy, Séjourné, Vialard, Amari, Trounev, and Peyré \(2019\)](#).

There is also a statistical interpretation. For quadratic cost, set

$$W_\sigma^2(\mu, \nu) = \inf_{\gamma \in \Pi(\mu, \nu)} \left\{ \frac{1}{2} \int \|\mathbf{x} - \mathbf{y}\|^2 d\gamma + \sigma^2 I(\gamma) \right\}, \quad (3.10)$$

where $I(\gamma) = \text{KL}(\gamma \mid \mu \otimes \nu)$.

Suppose that latent observations $X_1, \dots, X_n$ are independent with law $\mathbb{P}$, but only

$$Y_i = X_i + \varepsilon_i \quad \text{and} \quad \varepsilon_i \sim \mathcal{N}(0, \sigma^2 \mathbf{I})$$

are observed. Estimating $\mathbb{P}$ from the convolution of $\mathbb{P}$ with Gaussian noise is a **deconvolution problem**.

Theorem 3.8 Let $\mathcal{S}$ be a model class closed under domination: if $\mathbb{Q} \in \mathcal{S}$ and $\mathbb{M} \ll \mathbb{Q}$, then $\mathbb{M} \in \mathcal{S}$. The **nonparametric maximum likelihood estimator** for the deconvolution problem satisfies

$$\hat{\mathbb{P}} = \operatorname{argmin}_{\mathbb{Q} \in \mathcal{S}} W_\sigma^2 \left(\mathbb{Q}, \frac{1}{n} \sum_{i=1}^n \delta_{Y_i} \right).$$

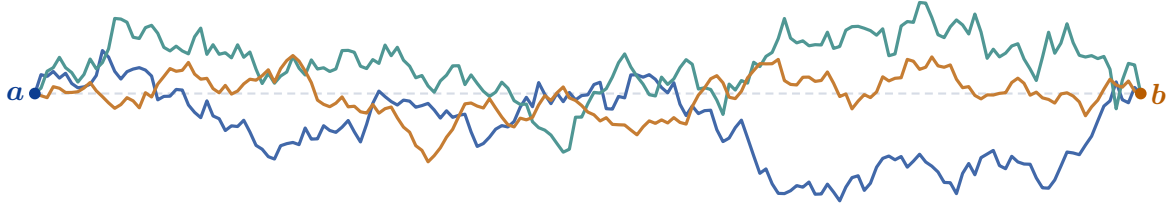
Thus entropic transport is not merely a numerical smoothing device: its mutual information penalty encodes a **latent variable noise model**. The result is due to [Rigollet and Weed \(2018\)](#).

Schrödinger Bridges and Large Deviations

Consider a cloud of independent particles evolving according to a reference stochastic law $\mathbb{M}$ on a path space Ω . Given an initial distribution μ_0 and an observed terminal distribution μ_1 , the **Schrödinger problem** asks for the most likely path law compatible with these endpoints.

For a single Brownian particle conditioned to travel from $\mathbf{a}$ to $\mathbf{b}$, the **Brownian bridge** fluctuates around the straight segment. Its continuous sample paths remain rough at every scale and have no classical velocity. The three realizations in [Figure 6](#) illustrate this irregularity. As the volatility vanishes, the random path collapses to the constant velocity path from [\(2.6\)](#).

More formally, take $\Omega = C([0, 1], \mathbb{R}^d)$ and let $X_t(\omega) = \omega(t)$ be the **canonical process**. Under **Wiener measure**, (X_t) is **Brownian motion**; under another probability on Ω , the same coordinate maps describe a different process.



Brownian bridge realizations

FIGURE 6

Definition 3.9 Let $\mathbb{M}$ be a reference probability on path space. The **dynamic Schrödinger problem** is

$$\inf_{\substack{\mathbb{P}: \mathbb{P}_0 = \mu_0 \\ \mathbb{P}_1 = \mu_1}} \text{KL}(\mathbb{P} \mid \mathbb{M}).$$

If $\mathbb{M}_{01}$ is the joint endpoint law, the **static Schrödinger problem** is

$$\inf_{\gamma \in \Pi(\mu_0, \mu_1)} \text{KL}(\gamma \mid \mathbb{M}_{01}).$$

Let $\mathbb{M}^{xy}$ be the **reference bridge** from x to y , and disintegrate a candidate path law as $\mathbb{P}^{xy} d\mathbb{P}_{01}(x, y)$. The **chain rule for relative entropy** gives

$$\text{KL}(\mathbb{P} \mid \mathbb{M}) = \text{KL}(\mathbb{P}_{01} \mid \mathbb{M}_{01}) + \int \text{KL}(\mathbb{P}^{xy} \mid \mathbb{M}^{xy}) d\mathbb{P}_{01}(x, y). \quad (3.11)$$

The second term is minimized by using the reference bridges themselves. Hence the static and dynamic problems are equivalent; the static optimizer chooses endpoint pairs, and the reference fills in the paths.

Suppose that, relative to a measure η on $X \times X$,

$$d\mathbb{M}_{01}^\varepsilon(x, y) \asymp Z(\varepsilon) \exp\left(-\frac{c(x, y)}{\varepsilon}\right) d\eta(x, y).$$

Writing $d\gamma = \rho d\eta$ gives

$$\varepsilon \text{KL}(\gamma \mid \mathbb{M}_{01}^\varepsilon) = \int c d\gamma + \varepsilon \int \rho \log \rho d\eta + \text{a constant}.$$

Thus the static Schrödinger problem is an entropically regularized transport problem.

Definition 3.10 Let $(\mathbb{P}^\varepsilon)_{\varepsilon>0}$ be probabilities on X , and let $I: X \rightarrow [0, \infty]$ be lower semi-continuous. The family satisfies a **large deviation principle** with **rate function** I if, for regular enough Borel sets A ,

$$\lim_{\varepsilon \downarrow 0} \varepsilon \log \mathbb{P}^\varepsilon(A) = - \inf_{x \in A} I(x).$$

Equivalently at the heuristic exponential scale,

$$\mathbb{P}^\varepsilon(A) \asymp \exp \left(-\frac{1}{\varepsilon} \inf_{x \in A} I(x) \right).$$

The rigorous definition uses separate lower and upper bounds for open and closed sets.

Theorem 3.11 (Sanov's theorem) Let X be finite and fix $\nu \in \mathcal{P}(X)$. If $\mathbb{P}_n$ is the law of the empirical measure

$$\frac{1}{n} \sum_{i=1}^n \delta_{X_i}, \quad X_i \text{ independent with law } \nu,$$

then $(\mathbb{P}_n)$ satisfies a large deviation principle with speed n and rate

$$I(\mathbf{p}) = \text{KL}(\mathbf{p} \mid \nu).$$

Theorem 3.12 (Schilder's theorem) Let $\mathbf{B}$ be standard Brownian motion in $\mathbb{R}^d$, and let $\mathbb{P}_k$ be the law of $\mathbf{B}/\sqrt{k}$ on $C([0, 1], \mathbb{R}^d)$. Then $(\mathbb{P}_k)$ satisfies a large deviation principle with rate

$$I(\omega) = \begin{cases} \frac{1}{2} \int_0^1 \|\dot{\omega}_t\|^2 dt, & \omega_0 = \mathbf{0} \text{ and } \omega \text{ is absolutely continuous,} \\ \infty, & \text{otherwise.} \end{cases}$$

The rate in [Theorem 3.12](#) is exactly the **kinetic action** in (2.6). This identifies quadratic transport as the **zero noise limit** of the Brownian Schrödinger bridge.

Theorem 3.13 For each k , let $\rho_k(x, \cdot)$ be a transition kernel and set

$$d\mathbb{M}_{01}^k(x, y) = d\mu_0(x) \rho_k(x, dy).$$

Suppose that, for each x , the kernels satisfy a large deviation principle with rate $c(x, \cdot)$ and speed k . Under the required tightness conditions,

$$\inf_{\gamma \in \Pi(\mu_0, \mu_1)} \frac{1}{k} \text{KL}(\gamma \mid \mathbb{M}_{01}^k) \longrightarrow \mathcal{T}_c(\mu_0, \mu_1).$$

Every cluster point of the Schrödinger optimizers is optimal for the Monge–Kantorovich problem.

For Brownian kernels, the **entropic interpolation** converges to the displacement interpolation in [Definition 2.15](#).

A Dirichlet Schrödinger Problem

The cost in (2.3) also arises as a large deviation rate. Recall that a random variable G has a gamma distribution with shape α and rate β if

$$f_G(y) = \frac{\beta^\alpha}{\Gamma(\alpha)} y^{\alpha-1} e^{-\beta y}, \quad y > 0.$$

If G_i are independent with $G_i \sim \text{Gamma}(\alpha_i, 1)$ and

$$D_i = \frac{G_i}{\sum_{j=1}^n G_j},$$

then $\mathbf{D} \sim \text{Dir}(\alpha_1, \dots, \alpha_n)$ has a **Dirichlet distribution**.

Use the multiplicative simplex operation to set $\mathbf{Q} = \mathbf{p} \oplus \mathbf{D}$.

Lemma 3.14 The conditional density of $\mathbf{Q}$ given $\mathbf{p}$ is

$$f(\mathbf{q} \mid \mathbf{p}) = \frac{\Gamma(\sum_j \alpha_j)}{\prod_j p_j \Gamma(\alpha_j)} \prod_{i=1}^n \left(\frac{q_i}{p_i} \right)^{\alpha_i-1} \left(\sum_{j=1}^n \frac{q_j}{p_j} \right)^{-\sum_j \alpha_j}.$$

Proposition 3.15 For each $\lambda > 0$, let $\alpha_i = \lambda/n$. Then

$$\lim_{\lambda \rightarrow \infty} -\frac{1}{\lambda} \log f_\lambda(\mathbf{q} \mid \mathbf{p}) = c(\mathbf{p}, \mathbf{q})$$

up to terms independent of $\mathbf{q}$. Thus the **Dirichlet kernels** satisfy a large deviation principle with the Dirichlet transport cost as rate function.

The Dirichlet distribution therefore plays the role that the normal distribution plays in quadratic transport.

To obtain a particle interpretation, draw

$$\mathbf{p}^{(1)}, \dots, \mathbf{p}^{(N)} \sim \mu_0 \quad \text{and} \quad \mathbf{q}^{(1)}, \dots, \mathbf{q}^{(N)} \sim \mu_1.$$

Independently draw $\mathbf{D}^{(j)} \sim \text{Dir}(\lambda/n, \dots, \lambda/n)$ and set

$$\tilde{\mathbf{q}}^{(j)} = \mathbf{p}^{(j)} \oplus \mathbf{D}^{(j)}.$$

Conditioning the empirical law of the $\tilde{\mathbf{q}}^{(j)}$ to equal the empirical target law produces a random coupling $\tilde{\gamma}_N$.

Theorem 3.16 Under the technical regularity assumptions, if λ_N is of order $N^{2/n}$, then $\tilde{\gamma}_N$ converges, with quantitative rates, to the solution of the Dirichlet Monge–Kantorovich problem.

The construction extends to other costs whenever an appropriate large deviation kernel is available. In this limit the c -divergence of [Definition 2.11](#) appears as the local cost of a mismatched endpoint pair.

Finally, a Schrödinger bridge also admits an **Eulerian form**. If

$$\partial_t \rho_t + \nabla \cdot (\rho_t \mathbf{v}) = 0, \quad \rho_0 = \mu, \quad \text{and} \quad \rho_1 = \nu,$$

then, under the Brownian reference with noise parameter ε , the bridge minimizes

$$\frac{1}{2} \int_0^1 \int_{\mathbb{R}^d} \left(\|\mathbf{v}(t, \mathbf{x})\|^2 + \frac{\varepsilon^2}{4} \|\nabla \log \rho_t(\mathbf{x})\|^2 \right) d\rho_t(\mathbf{x}) dt. \quad (3.12)$$

The first term is **kinetic energy**. The second is the time integral of the Fisher information and penalizes sharply concentrated intermediate densities.

Lecture 4 — Friday, October 04, 2019

3. Riemannian and Information Geometry

Manifolds, Tangent Spaces, and Geodesics

A **smooth n -dimensional manifold** M is locally modelled on $\mathbb{R}^n$. A **chart** identifies a neighbourhood $U \subseteq M$ with coordinates in $\mathbb{R}^n$, and the **changes of coordinates** are diffeomorphisms. The same language can be used heuristically for spaces of probability measures. When the state space has $n + 1$ points, the strictly positive distributions form the finite dimensional simplex

$$S_n = \left\{ \mathbf{p} = (p_0, \dots, p_n) : p_i > 0, \sum_{i=0}^n p_i = 1 \right\}.$$

For a continuous state space, choosing a topology on a space of densities is part of the analysis rather than a formality.

The local geometry is illustrated in Figure 7 by two overlapping **coordinate charts**. Their overlap is encoded by a smooth change of coordinates.

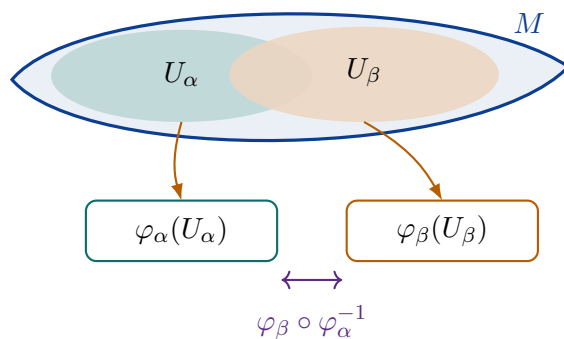


FIGURE 7

The **tangent space** $T_p M$ consists of **velocity vectors** of smooth curves through p . If $\gamma(0) = p$, its velocity acts on a smooth function by

$$\dot{\gamma}(f) = \left. \frac{d}{dt}(f \circ \gamma)(t) \right|_{t=0}.$$

Definition 4.1 A **Riemannian metric** is a smoothly varying family of inner products

$$g_p: T_p M \times T_p M \rightarrow \mathbb{R}.$$

In a basis (e_i) of $T_p M$,

$$g_p(\mathbf{u}, \mathbf{v}) = \sum_{i,j} g_{ij}(p) u^i v^j.$$

The length of a smooth curve $\gamma: [0, 1] \rightarrow M$ is

$$\ell(\gamma) = \int_0^1 \sqrt{g_{\gamma(t)}(\dot{\gamma}(t), \dot{\gamma}(t))} dt,$$

and the **Riemannian distance** is the infimum of this length over curves joining the endpoints. A length minimizing curve is a **geodesic**. Equivalently, we may minimize the energy

$$\int_0^1 g_{\gamma(t)}(\dot{\gamma}(t), \dot{\gamma}(t)) dt.$$

In coordinates, an **affinely parameterized geodesic** solves

$$\ddot{\gamma}^k(t) + \sum_{i,j} \Gamma_{ij}^k(\gamma(t)) \dot{\gamma}^i(t) \dot{\gamma}^j(t) = 0.$$

The coefficients Γ_{ij}^k are the **Christoffel symbols** of the metric.

Otto's Wasserstein Geometry

The [Benamou–Brenier formula](#) in [Theorem 2.21](#) suggests a **weak Riemannian structure** on the space of smooth positive densities. Let M be a compact Riemannian manifold and set

$$\mathcal{P}_+^\infty(M) = \left\{ \rho \in C^\infty(M) : \rho > 0, \int_M \rho d\text{vol} = 1 \right\}.$$

A **tangent vector** at ρ is a smooth function σ of zero integral. If

$$\sigma = -\text{div}(\rho \nabla \varphi),$$

define **Otto's metric** by

$$g_\rho^W(\sigma, \sigma) = \int_M g_\rho^M(\nabla \varphi, \nabla \varphi) \rho(\mathbf{x}) d\text{vol}(\mathbf{x}). \quad (4.1)$$

The continuity equation explains the definition: if $t \mapsto \rho_t$ has velocity field $\nabla\varphi$ at a given time, then its tangent is $\dot{\rho}_t = -\operatorname{div}(\rho_t \nabla\varphi)$.

In this notation, [Theorem 2.21](#) becomes the **formal Riemannian energy representation**

$$W_2(\mu_0, \mu_1)^2 = \inf_{\substack{\rho_0=\mu_0 \\ \rho_1=\mu_1}} \int_0^1 g_{\rho_t}^W(\dot{\rho}_t, \dot{\rho}_t) dt. \quad (4.2)$$

This interpretation, associated with [Otto \(2001\)](#), is formal but predicts the correct differential equations and gradient flows.

Fisher–Rao Geometry and Sufficiency

Consider a smooth parametric family of densities $p(\mathbf{x}; \boldsymbol{\theta})$, with $\boldsymbol{\theta} \in \Theta \subseteq \mathbb{R}^d$, and denote the corresponding probability measures by $\mathbb{P}_{\boldsymbol{\theta}}$.

Definition 4.2 The **Fisher information matrix** is

$$G_{ij}(\boldsymbol{\theta}) = \mathbb{E}_{\boldsymbol{\theta}} \left[\frac{\partial}{\partial \theta_i} \log p(X; \boldsymbol{\theta}) \frac{\partial}{\partial \theta_j} \log p(X; \boldsymbol{\theta}) \right]. \quad (4.3)$$

The Cramér–Rao inequality states, under the usual fixed-support differentiation conditions and when $\mathbf{G}(\boldsymbol{\theta})$ is nonsingular, that an unbiased estimator $\hat{\boldsymbol{\theta}}$ satisfies

$$\operatorname{Cov}_{\boldsymbol{\theta}}(\hat{\boldsymbol{\theta}}) \succeq \mathbf{G}(\boldsymbol{\theta})^{-1}.$$

Rao’s geometric insight was to regard $\mathbf{G}$ as a Riemannian metric rather than only an information bound.

The same metric is the second-order part of relative entropy.

Proposition 4.3 Suppose that the log density is three times differentiable in a neighbourhood of $\boldsymbol{\theta}$, with derivatives dominated sufficiently to interchange differentiation and integration. If $\boldsymbol{\theta}' = \boldsymbol{\theta} + \Delta\boldsymbol{\theta}$, then

$$\operatorname{KL}(\mathbb{P}_{\boldsymbol{\theta}'} \mid \mathbb{P}_{\boldsymbol{\theta}}) = \frac{1}{2} \Delta\boldsymbol{\theta}^\top \mathbf{G}(\boldsymbol{\theta}) \Delta\boldsymbol{\theta} + O(\|\Delta\boldsymbol{\theta}\|^3).$$

The reverse relative entropy has the same quadratic term.

To see why the metric is intrinsic, let μ be a reference measure and consider the **formal tangent**

space at a positive density ρ . For a tangent σ with integral zero, set

$$g_\rho^F(\sigma, \sigma) = \int_X \frac{|\sigma|^2}{\rho} d\mu. \quad (4.4)$$

Write $p_\theta = d\mathbb{P}_\theta / d\mu$. If $\mathbf{a} \in T_\theta \Theta$, its **pushforward tangent** under $\theta \mapsto p_\theta$ is

$$\sigma = \sum_i a_i \frac{\partial p_\theta}{\partial \theta_i}.$$

Substitution into (4.4) gives

$$g_{p_\theta}^F(\sigma, \sigma) = \sum_{i,j} a_i a_j \int_X \frac{1}{p_\theta} \frac{\partial p_\theta}{\partial \theta_i} \frac{\partial p_\theta}{\partial \theta_j} d\mu = \mathbf{a}^\top \mathbf{G}(\theta) \mathbf{a}.$$

Thus the Fisher information metric is the **pullback** of an **intrinsic metric** on densities.

On the **cone of positive measures**, its squared distance is represented dynamically by

$$d_{FR}(\rho_0, \rho_1)^2 = \inf_{\partial_t \rho = \zeta} \int_0^1 \int_X \frac{|\zeta(t, x)|^2}{\rho(t, x)} d\mu(x) dt. \quad (4.5)$$

The change of variable $r = 2\sqrt{\rho}$ makes the metric Euclidean, so a **positive measure geodesic** is

$$\rho_t(x) = \left((1-t)\sqrt{\rho_0(x)} + t\sqrt{\rho_1(x)} \right)^2.$$

If total mass is constrained to equal one, the **square-root densities** lie on a sphere, and geodesics are **great circle arcs** on that sphere.

For a Riemannian manifold, the gradient of F is determined by

$$g_p(\text{grad } F(p), \mathbf{X}) = \mathbf{X}F(p).$$

On the simplex with the Fisher metric,

$$(\text{grad } F(\mathbf{x}))_i = x_i \left(\frac{\partial F}{\partial x_i} - \sum_j x_j \frac{\partial F}{\partial x_j} \right). \quad (4.6)$$

The negative gradient flow is the **replicator equation** from evolutionary game theory.

The principal reason to privilege the Fisher metric is its invariance under sufficient statistics.

Definition 4.4 For a family $(\mathbb{P}_\theta)$ and an observation X , a statistic $T(X)$ is a **sufficient statistic** if the conditional distribution

$$\mathbb{P}_\theta(X \in \cdot \mid T(X))$$

does not depend on θ .

On a finite state space $I = \{0, \dots, n\}$, a statistic determines a partition

$$T: I \rightarrow \{A_0, \dots, A_m\}.$$

The induced **coarse graining map** $f: S_n \rightarrow S_m$ is

$$f(\mathbf{p}) = \mathbf{q} \quad \text{and} \quad q_j = \sum_{i \in A_j} p_i.$$

Sufficiency for a model $M \subseteq S_n$ is equivalent to the factorization

$$p_i = r_{ij} q_j, \quad i \in A_j \quad \text{and} \quad \sum_{i \in A_j} r_{ij} = 1,$$

where the conditional probabilities r_{ij} are fixed over the model. This defines the embedding $h: S_m \rightarrow S_n$ in Figure 8, with $f \circ h = \text{Id}$.

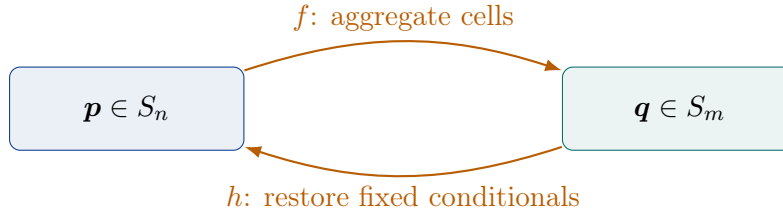


FIGURE 8

Theorem 4.5 Suppose that each simplex S_n , for $n \geq 2$, is equipped with a Riemannian metric $g^{(n)}$. Assume that every sufficient statistic embedding $h: S_m \rightarrow S_n$ of the preceding form is an isometry onto its image. Then

$$g^{(n)} = c g_F^{(n)}$$

for a constant $c > 0$ independent of n , where

$$g_{F,\mathbf{p}}^{(n)}(\mathbf{u}, \mathbf{v}) = \sum_{i=0}^n \frac{u_i v_i}{p_i}.$$

Proof. At the uniform distribution

$$\bar{\mathbf{p}} = \left(\frac{1}{n+1}, \dots, \frac{1}{n+1} \right),$$

every permutation of the coordinates is an isometry. Permutation invariance forces the ambient coordinate representation of the metric to have the form

$$g_{ij}^{(n)}(\bar{\mathbf{p}}) = C_n \delta_{ij} + D_n.$$

Tangent vectors satisfy $\sum_i u_i = \sum_i v_i = 0$, so the term involving D_n vanishes. We may therefore write

$$g_{\bar{\mathbf{p}}}^{(n)}(\mathbf{u}, \mathbf{v}) = C_n \sum_i u_i v_i.$$

Fix m and a rational point $\mathbf{q} \in S_m$ of the form

$$q_j = \frac{k_j}{n+1}, \quad k_j \in \mathbb{N} \quad \text{and} \quad \sum_{j=0}^m k_j = n+1.$$

Choose a partition with $|A_j| = k_j$ and the uniform conditional probabilities $r_{ij} = 1/k_j$ for $i \in A_j$. Then $h(\mathbf{q}) = \bar{\mathbf{p}}$, and the differential satisfies $(dh(\mathbf{u}))_i = u_j/k_j$ for $i \in A_j$. The isometry assumption gives

$$\begin{aligned} g_{\mathbf{q}}^{(m)}(\mathbf{u}, \mathbf{v}) &= C_n \sum_{j=0}^m \sum_{i \in A_j} \frac{u_j}{k_j} \frac{v_j}{k_j} \\ &= \frac{C_n}{n+1} \sum_{j=0}^m \frac{u_j v_j}{q_j}. \end{aligned}$$

The left side is independent of the particular denominator $n+1$, so $C_n/(n+1) = c$ for a constant $c > 0$. Rational points are dense in S_m ; continuity therefore extends the formula to every point. $\square$

Relative entropy is itself invariant under sufficient statistic transformations. More generally this

holds for every f -divergence. Relative entropy is distinguished further as the only divergence on a finite simplex that is simultaneously an f -divergence and a Bregman divergence.

For one dimensional Gaussian distributions, the Fisher metric produces a familiar geometry. In the coordinates (m, σ) , its line element is

$$ds_F^2 = \frac{(dm)^2 + 2(d\sigma)^2}{\sigma^2}.$$

After setting $s = \sqrt{2}\sigma$, this becomes

$$ds_F^2 = 2 \frac{(dm)^2 + (ds)^2}{s^2}.$$

Thus the rescaled parameter half-plane (m, s) carries a constant multiple of the **Poincaré metric**, whose geodesics are the vertical lines and boundary-orthogonal semicircles shown in [Figure 9](#).

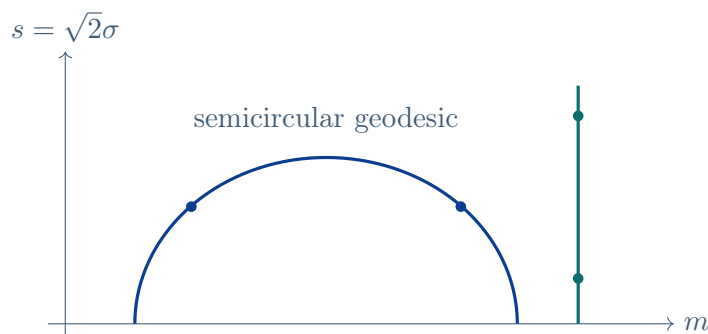


FIGURE 9

Wasserstein–Fisher–Rao Interpolation

The **Wasserstein action** transports mass using **momentum** $\omega = \rho v$, while the **Fisher–Rao action** creates or destroys mass through a **source** ζ :

$$\inf_{\partial_t \rho + \nabla \cdot \omega = 0} \int_0^1 \int_{\mathbb{R}^n} \frac{\|\omega\|^2}{\rho} d\mathbf{x} dt$$

and

$$\inf_{\partial_t \rho = \zeta} \int_0^1 \int_{\mathbb{R}^n} \frac{|\zeta|^2}{\rho} d\mathbf{x} dt.$$

Combining the continuity laws gives **unbalanced transport**.

Definition 4.6 For $\kappa > 0$ and nonnegative finite measures ρ_0 and ρ_1 , the **Wasserstein–Fisher–Rao distance** is defined by

$$WF_\kappa(\rho_0, \rho_1)^2 = \inf_{\rho, \omega, \zeta} \int_0^1 \left[\int_{\mathbb{R}^n} \frac{\|\omega(t, \mathbf{x})\|^2}{\rho(t, \mathbf{x})} d\mathbf{x} + \kappa^2 \int_{\mathbb{R}^n} \frac{|\zeta(t, \mathbf{x})|^2}{\rho(t, \mathbf{x})} d\mathbf{x} \right] dt, \quad (4.7)$$

subject to

$$\partial_t \rho + \nabla \cdot \omega = \zeta, \quad \rho(0) = \rho_0, \quad \text{and} \quad \rho(1) = \rho_1.$$

The action uses the lower-semicontinuous convention that a quotient with $\rho = 0$ is zero when its numerator is zero and is $+\infty$ otherwise. For general measures, (4.7) denotes the corresponding measure-theoretic relaxation.

The parameter κ controls the cost of changing mass relative to moving it. Large κ favours classical transport; small κ permits mass to disappear at one location and reappear elsewhere.

Partial transport provides a related **static model**. If $\Pi_{\leq}(\rho_0, \rho_1)$ denotes measures whose marginals are dominated by ρ_0 and ρ_1 , one minimizes

$$\inf_{\substack{\gamma \in \Pi_{\leq}(\rho_0, \rho_1) \\ \gamma(X \times Y) = m}} \int c d\gamma.$$

Its dynamic formulation replaces the quadratic source cost in (4.7) by an **L^1 source penalty**. More generally, for $p, q > 1$ one considers

$$\inf \int_0^1 \left[\frac{1}{p} \int \frac{\|\omega\|^p}{\rho^{p-1}} + \frac{\kappa^q}{q} \int \frac{|\zeta|^q}{\rho^{q-1}} \right] dt.$$

The existence proof for (4.7) is based on convex duality.

Theorem 4.7 (Fenchel–Rockafellar duality) Let E and F be locally convex Hausdorff spaces, let $A: E \rightarrow F$ be continuous and linear, and let f and g be proper lower semicontinuous convex functions. If f is finite at some $x \in E$ at which g is continuous at Ax , and if

the primal value below is finite, then

$$\inf_{x \in E} \{f(x) + g(Ax)\} = \max_{y^* \in F^*} \{-f^*(-A^*y^*) - g^*(y^*)\}.$$

The dual maximum is attained. Primal and dual optimizers satisfy the corresponding sub-gradient conditions.

Theorem 4.8 The action in (4.7) admits an optimizer. Moreover:

1. WF_κ is a metric on the space of finite nonnegative measures;
2. its minimizing interpolation is a constant speed geodesic,

$$WF_\kappa(\rho_s, \rho_t) = |t - s|WF_\kappa(\rho_0, \rho_1);$$

3. for $\alpha > 0$,

$$WF_\kappa(\alpha\rho_0, \alpha\rho_1) = \sqrt{\alpha}WF_\kappa(\rho_0, \rho_1).$$

Numerically, the interpolation can combine translation, deformation, disappearance, and creation. Two nearby modes primarily move; sufficiently distant modes may shrink at the source and grow at the target, as in Figure 10. Shape interpolation displays the same distinction: Wasserstein motion deforms the outline continuously, Fisher–Rao motion adjusts its intensity, and the combined metric selects the cheaper mechanism.

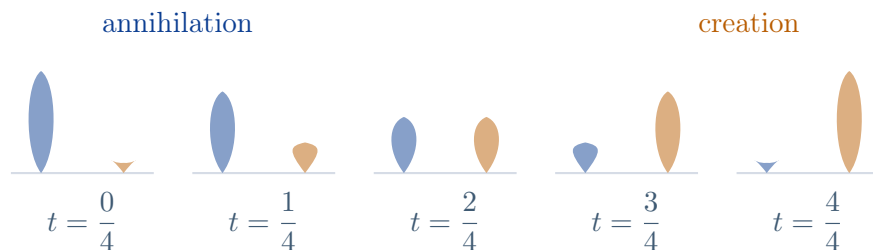


FIGURE 10

These results and numerical schemes are developed by Chizat, Peyré, Schmitzer, and Vialard (2018).

Natural Gradients and Wasserstein Statistical Manifolds

Euclidean gradient descent iterates

$$\boldsymbol{\theta}^{(k+1)} = \boldsymbol{\theta}^{(k)} - \eta \nabla F(\boldsymbol{\theta}^{(k)}).$$

It can be viewed as minimizing the **local model**

$$F(\boldsymbol{\theta}) + \frac{1}{2\eta} \left\| \boldsymbol{\theta} - \boldsymbol{\theta}^{(k)} \right\|^2.$$

The Euclidean penalty is coordinate dependent. On a Riemannian model with metric matrix $\mathbf{G}(\boldsymbol{\theta})$, the gradient satisfies

$$\text{grad } F(\boldsymbol{\theta}) = \mathbf{G}(\boldsymbol{\theta})^{-1} \nabla F(\boldsymbol{\theta}).$$

This gives the **natural gradient update**

$$\boldsymbol{\theta}^{(k+1)} = \boldsymbol{\theta}^{(k)} - \eta \mathbf{G}(\boldsymbol{\theta}^{(k)})^{-1} \nabla F(\boldsymbol{\theta}^{(k)}). \quad (4.8)$$

With the Fisher metric, the update searches inside a **local relative entropy ball**. Empirical neural network examples show that this **geometric preconditioning** can reduce the number of epochs, though computing and inverting $\mathbf{G}$ introduces its own cost.

The same construction can use the Wasserstein metric. Let

$$\mathcal{M} = \{\rho(\cdot, \boldsymbol{\theta}) : \boldsymbol{\theta} \in \Theta\} \subseteq \mathcal{P}_2(\mathbb{R}^d).$$

Pulling the Wasserstein metric back to Θ gives a **Wasserstein statistical manifold**. Its **intrinsic geodesics** are constrained to remain in $\mathcal{M}$ and therefore need not equal the ambient displacement interpolation. The two agree when $\mathcal{M}$ is **geodesically closed**, as happens for Gaussian families.

In one state space dimension the pullback is explicit.

Proposition 4.9 Let

$$F(x, \boldsymbol{\theta}) = \int_{-\infty}^x \rho(t, \boldsymbol{\theta}) \, dt.$$

The Wasserstein metric matrix is

$$\mathbf{G}_W(\boldsymbol{\theta}) = \int_{\mathbb{R}} \frac{1}{\rho(x, \boldsymbol{\theta})} \nabla_{\boldsymbol{\theta}} F(x, \boldsymbol{\theta}) \nabla_{\boldsymbol{\theta}} F(x, \boldsymbol{\theta})^{\top} \, dx.$$

For comparison, the Fisher metric is

$$\mathbf{G}_F(\boldsymbol{\theta}) = \int_{\mathbb{R}} \frac{1}{\rho(x, \boldsymbol{\theta})} \nabla_{\boldsymbol{\theta}} \rho(x, \boldsymbol{\theta}) \nabla_{\boldsymbol{\theta}} \rho(x, \boldsymbol{\theta})^\top dx.$$

For a mixture of two normal densities, Wasserstein geodesics within the model move the mixture components smoothly, whereas a straight line in the density coordinates may create intermediate multimodal shapes. In parameter estimation, we may minimize

$$\min_{\boldsymbol{\theta}} W_2 \left(\rho(\cdot, \boldsymbol{\theta}), \frac{1}{N} \sum_{j=1}^N \delta_{x_j} \right)^2$$

using the **Wasserstein natural gradient**. The one dimensional monotone coupling makes the loss explicit, but scalability remains the main computational difficulty.

Entropy Relaxations and Pseudo-Riemannian Connections

For a finite state space, define the **Cuturi function**

$$C_\lambda(\mathbf{p}, \mathbf{q}) = \frac{1}{1 + \lambda} (\langle \mathbf{M}, \mathbf{P}_\lambda^* \rangle - \lambda H(\mathbf{P}_\lambda^*)), \quad (4.9)$$

where $\mathbf{P}_\lambda^*$ is the entropically regularized optimal coupling with marginals $\mathbf{p}$ and $\mathbf{q}$. By [Theorem 3.4](#),

$$P_{ij}^* = u_i v_j \exp(-M_{ij}/\lambda).$$

Equivalently, after absorbing normalization into potentials,

$$P_{ij}^* \propto \exp \left(\frac{1 + \lambda}{\lambda} (\alpha_i + \beta_j) - \frac{1}{\lambda} M_{ij} \right).$$

Consequently, the family of all regularized optimal couplings obtained by varying the marginals is an exponential family, and the convex conjugate of its cumulant generating function is C_λ .

For a fixed reference distribution $\mathbf{r}$, define

$$D_{\mathbf{r}, \lambda}[\mathbf{p} : \mathbf{q}] = \frac{\lambda}{1 + \lambda} \text{KL}(\mathbf{P}_\lambda^*(\mathbf{r}, \mathbf{p}) \mid \mathbf{P}_\lambda^*(\mathbf{r}, \mathbf{q})).$$

When $\mathbf{r} = \mathbf{p}$, this has the following **representation resembling a Bregman divergence**:

$$D_{\mathbf{p}, \lambda}[\mathbf{p} : \mathbf{q}] = C_\lambda(\mathbf{p}, \mathbf{p}) - C_\lambda(\mathbf{p}, \mathbf{q}) - \nabla_{\mathbf{q}} C_\lambda(\mathbf{p}, \mathbf{q}) \cdot (\mathbf{p} - \mathbf{q}),$$

and

$$\lim_{\lambda \rightarrow \infty} D_{r,\lambda}[\mathbf{p} : \mathbf{q}] = \text{KL}(\mathbf{p} \mid \mathbf{q}).$$

This construction provides one interpolation between optimal transport and information geometry, but it does not by itself give a single geometry encompassing every such interpolation.

A different connection begins with the **cross-difference** of a cost $c: M \times M' \rightarrow \mathbb{R}$:

$$\delta(p, q', p', q'') = c(p, q'') + c(p', q') - c(p, q') - c(p', q''). \quad (4.10)$$

If (p, q') and (p', q'') lie on an optimal graph, then $\delta \geq 0$ by cyclical monotonicity. Taylor expansion of the cross-difference yields a **pseudo-Riemannian metric** on $M \times M'$. In **local coordinates** (ξ, η') it has block form

$$\mathbf{h} = \frac{1}{2} \begin{pmatrix} \mathbf{0} & -\overline{D}Dc \\ -D\overline{D}c & \mathbf{0} \end{pmatrix}. \quad (4.11)$$

The **Ma–Trudinger–Wang tensor** is an unnormalized **cross-sectional curvature** of this metric.

Restrict now to the optimal graph G and the c -divergence from [Definition 2.11](#). For $x = (p, q)$ and $x' = (p', q')$ in G , write

$$\delta(x, x') = \delta(p, q, p', q').$$

Theorem 4.10 For $x, x' \in G$,

$$\delta(x, x') = D[x : x'] + D[x' : x].$$

Thus the cross-difference is the **symmetrization** of the c -divergence. This identity connects optimality, **divergence geometry**, and the **pseudo-Riemannian curvature framework**.

Lecture 5 — Friday, October 11, 2019

4. Stochastic Control and Extensions

Hamilton–Jacobi Duality

For quadratic transport, the static and dynamic primal problems are

$$\inf_{\gamma \in \Pi(\mu, \nu)} \frac{1}{2} \int \|\mathbf{x} - \mathbf{y}\|^2 d\gamma(\mathbf{x}, \mathbf{y})$$

and, by [Theorem 2.21](#),

$$\inf_{\rho, \mathbf{v}} \int_0^1 \int_{\mathbb{R}^n} \frac{1}{2} \|\mathbf{v}_t(\mathbf{x})\|^2 d\rho_t(\mathbf{x}) dt.$$

The dynamic formulation leads to a **Hamilton–Jacobi description** of the Kantorovich potentials.

If $(\rho, \mathbf{v})$ satisfies the continuity equation, then every smooth compactly supported test function φ satisfies

$$\int_0^1 \int_{\mathbb{R}^n} (\partial_t \varphi + \nabla \varphi \cdot \mathbf{v}) d\rho_t dt = \int_{\mathbb{R}^n} \varphi(1, \mathbf{y}) d\nu(\mathbf{y}) - \int_{\mathbb{R}^n} \varphi(0, \mathbf{x}) d\mu(\mathbf{x}). \quad (5.1)$$

Introduce the momentum $\boldsymbol{\omega} = \rho \mathbf{v}$ and enforce (5.1) by a Lagrange multiplier. Formally interchanging the infimum and supremum produces the pointwise minimization

$$\inf_{\rho \geq 0, \boldsymbol{\omega}} \left\{ \frac{\|\boldsymbol{\omega}\|^2}{2\rho} - (\partial_t \varphi)\rho - \nabla \varphi \cdot \boldsymbol{\omega} \right\}.$$

For fixed ρ , the minimizer is $\boldsymbol{\omega} = \rho \nabla \varphi$. The pointwise infimum over $\rho \geq 0$ is finite only for the dual subsolutions

$$\partial_t \varphi + \frac{1}{2} \|\nabla \varphi\|^2 \leq 0.$$

On the support of an optimal density, complementary slackness gives equality:

$$\partial_t \varphi + \frac{1}{2} \|\nabla \varphi\|^2 = 0. \quad (5.2)$$

At optimality, $-\varphi(0, \cdot)$ and $\varphi(1, \cdot)$ are Kantorovich potentials. Thus the static dual pair can be joined by a **Hamilton–Jacobi evolution**.

Stochastic Control and Girsanov’s Theorem

Let $\varepsilon > 0$ be a noise parameter. Consider continuous semimartingales of the form

$$\mathbf{X}_t = \mathbf{X}_0 + \int_0^t \boldsymbol{\beta}_{\mathbf{X}}(s, \mathbf{X}) ds + \sqrt{\varepsilon} \mathbf{W}_{\mathbf{X}}(t), \quad (5.3)$$

where $\mathbf{W}_\mathbf{X}$ is Brownian motion adapted to the **filtration** of $\mathbf{X}$. Given **endpoint laws** $\mathbb{P}_0$ and $\mathbb{P}_1$, define

$$V_\varepsilon(\mathbb{P}_0, \mathbb{P}_1) = \inf \mathbb{E} \left[\int_0^1 \frac{1}{2} \|\beta_\mathbf{X}(t, \mathbf{X})\|^2 dt \right], \quad (5.4)$$

where the infimum is over processes satisfying (5.3), $\mathbf{X}_0 \sim \mathbb{P}_0$, and $\mathbf{X}_1 \sim \mathbb{P}_1$.

Unlike a **standard control problem**, both endpoint distributions are constrained. A more customary **controlled diffusion** begins from a fixed state,

$$\mathbf{X}_t^\alpha = \mathbf{x}_0 + \int_0^t \mathbf{b}(s, \mathbf{X}_s^\alpha, \alpha_s) ds + \int_0^t \boldsymbol{\sigma}(s, \mathbf{X}_s^\alpha, \alpha_s) d\mathbf{W}_s,$$

and minimizes

$$J(\alpha) = \mathbb{E} \left[\int_0^T f(t, \mathbf{X}_t^\alpha) dt + g(\mathbf{X}_T^\alpha) \right].$$

The endpoint law constraint is what ties (5.4) to transport and Schrödinger bridges.

Suppose that $X \sim \mathcal{N}(0, \sigma^2)$ under $\mathbb{P}$. The following **change of measure** displays the basic mechanism:

$$Z = \exp \left(-\frac{\theta}{\sigma^2} X - \frac{\theta^2}{2\sigma^2} \right).$$

The density Z defines a probability $\mathbb{Q}$ by $d\mathbb{Q} = Z d\mathbb{P}$ under which $X \sim \mathcal{N}(-\theta, \sigma^2)$. [Girsanov's theorem](#) applies the same infinitesimal shift to Brownian increments.

Theorem 5.1 (Girsanov's theorem) Let $(\boldsymbol{\theta}_t)_{0 \leq t \leq T}$ be adapted.

Suppose that the **exponential martingale**

$$Z_T = \exp \left(-\int_0^T \boldsymbol{\theta}_t \cdot d\mathbf{W}_t - \frac{1}{2} \int_0^T \|\boldsymbol{\theta}_t\|^2 dt \right)$$

has expectation one. Define $\mathbb{Q}$ by $d\mathbb{Q} = Z_T d\mathbb{P}$. Then

$$\widetilde{\mathbf{W}}_t = \mathbf{W}_t + \int_0^t \boldsymbol{\theta}_s ds$$

is Brownian motion under $\mathbb{Q}$.

Apply [Theorem 5.1](#) on **canonical Wiener space** to

$$\mathbf{X}_t = \mathbf{X}_0 + \sqrt{\varepsilon} \mathbf{W}_t.$$

Choosing $\boldsymbol{\theta} = -\boldsymbol{\beta}/\sqrt{\varepsilon}$ turns the canonical process under the new measure into

$$\mathbf{X}_t = \mathbf{X}_0 + \int_0^t \boldsymbol{\beta}_s \, ds + \sqrt{\varepsilon} \widetilde{\mathbf{W}}_t.$$

The cost of this change of measure is relative entropy. Here $\mathbb{P}$ is the reference law and $\mathbb{Q}$ is the controlled law. The stochastic integral has mean zero under the controlled law, and

$$\text{KL}(\mathbb{Q} \mid \mathbb{P}) = \frac{1}{2\varepsilon} \mathbb{E}_{\mathbb{Q}} \left[\int_0^1 \|\boldsymbol{\beta}_t\|^2 \, dt \right]. \quad (5.5)$$

Equation (5.5) makes the equivalence between the **control energy** and the Schrödinger relative entropy transparent.

Noise adds diffusion to the **Hamilton–Jacobi equation**. The resulting **viscous Hamilton–Jacobi–Bellman equation** is

$$\partial_t \varphi + \frac{\varepsilon}{2} \Delta \varphi + \frac{1}{2} \|\nabla \varphi\|^2 = 0. \quad (5.6)$$

Theorem 5.2 Under the regularity assumptions for the **control problem with fixed endpoints**,

$$V_\varepsilon(\mathbb{P}_0, \mathbb{P}_1) = \sup_{\varphi} \left\{ \int \varphi(1, \mathbf{y}) \, d\mathbb{P}_1(\mathbf{y}) - \int \varphi(0, \mathbf{x}) \, d\mathbb{P}_0(\mathbf{x}) \right\},$$

where the supremum is over smooth subsolutions of

$$\partial_t \varphi + \frac{\varepsilon}{2} \Delta \varphi + \frac{1}{2} \|\nabla \varphi\|^2 \leq 0,$$

or over suitably interpreted viscosity subsolutions. Under conditions that permit solving the backward equation from arbitrary terminal data, the supremum may equivalently be parametrized by solutions of (5.6).

As $\varepsilon \downarrow 0$, (5.6) approaches (5.2), and the **stochastic dual** converges to **Kantorovich duality**. This formulation is due to Mikami and Thieullen (2006) and Mikami and Thieullen (2008).

Transport with Prior Dynamics

Quadratic transport can be generalized by imposing a **preferred linear dynamics**:

$$d\mathbf{x}(t) = \mathbf{A}(t)\mathbf{x}(t) \, dt + \mathbf{B}(t)\mathbf{u}(t) \, dt + \sqrt{\varepsilon}\mathbf{B}(t) \, d\mathbf{w}(t). \quad (5.7)$$

The state satisfies $\mathbf{x}(0) \sim \mu_0$ and $\mathbf{x}(1) \sim \mu_1$, and the objective is

$$\mathbb{E} \left[\int_0^1 \frac{1}{2} \|\mathbf{u}(t)\|^2 dt \right].$$

When $\mathbf{A} = \mathbf{0}$, $\mathbf{B} = \mathbf{I}$, and $\varepsilon = 0$, this is the [Benamou–Brenier problem](#). A nonzero $\mathbf{A}$ models an existing flow or force that the control should exploit rather than cancel unnecessarily.

Let $\Phi(t, s)$ be the transition matrix satisfying

$$\partial_t \Phi(t, s) = \mathbf{A}(t) \Phi(t, s) \quad \text{and} \quad \Phi(s, s) = \mathbf{I},$$

and define the **controllability Gramian**

$$\mathbf{M}(t_1, t_0) = \int_{t_0}^{t_1} \Phi(t_1, \tau) \mathbf{B}(\tau) \mathbf{B}(\tau)^\top \Phi(t_1, \tau)^\top d\tau. \quad (5.8)$$

The system is **controllable** when this matrix is invertible. Probabilistically, control can then steer the process from any initial point into any nonempty open target set with positive probability.

Set $\Phi_{10} = \Phi(1, 0)$ and $\mathbf{M}_{10} = \mathbf{M}(1, 0)$. In the deterministic case, the minimum energy required to travel from $\mathbf{x}$ to $\mathbf{y}$ is

$$c(\mathbf{x}, \mathbf{y}) = \frac{1}{2} (\mathbf{y} - \Phi_{10} \mathbf{x})^\top \mathbf{M}_{10}^{-1} (\mathbf{y} - \Phi_{10} \mathbf{x}). \quad (5.9)$$

The optimal control is

$$\mathbf{u}^*(t) = \mathbf{B}(t)^\top \Phi(1, t)^\top \mathbf{M}_{10}^{-1} (\mathbf{y} - \Phi_{10} \mathbf{x}),$$

and the associated state trajectory is

$$\mathbf{x}^*(t) = \Phi(t, 1) \mathbf{M}(1, t) \mathbf{M}_{10}^{-1} \Phi_{10} \mathbf{x} + \mathbf{M}(t, 0) \Phi(1, t)^\top \mathbf{M}_{10}^{-1} \mathbf{y}.$$

The change of variables

$$\hat{\mathbf{x}} = \mathbf{M}_{10}^{-1/2} \Phi_{10} \mathbf{x} \quad \text{and} \quad \hat{\mathbf{y}} = \mathbf{M}_{10}^{-1/2} \mathbf{y}$$

turns (5.9) into $\|\hat{\mathbf{x}} - \hat{\mathbf{y}}\|^2/2$. Therefore the problem reduces to quadratic transport between the transformed endpoint laws. If $\hat{T} = \nabla \phi$ is the corresponding Brenier map, then

$$T(\mathbf{x}) = \mathbf{M}_{10}^{1/2} \hat{T}(\mathbf{M}_{10}^{-1/2} \Phi_{10} \mathbf{x}). \quad (5.10)$$

The interpolating map is

$$T_t(\mathbf{x}) = \Phi(t, 1)\mathbf{M}(1, t)\mathbf{M}_{10}^{-1}\Phi_{10}\mathbf{x} + \mathbf{M}(t, 0)\Phi(1, t)^{\top}\mathbf{M}_{10}^{-1}T(\mathbf{x}),$$

which is generally not a straight line displacement interpolation.

The difference is visible even in one dimension. Under the prior drift $dx = -2x dt + u dt$, uncontrolled particles contract toward the origin. The optimal controlled paths in Figure 11 curve so that they cooperate with this drift, whereas **McCann interpolation** ignores it.

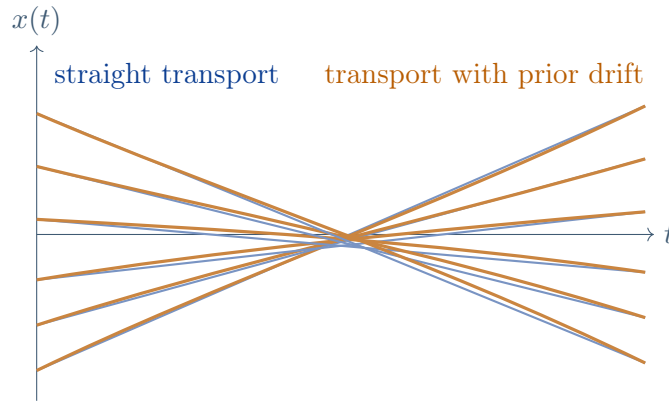


FIGURE 11

The same transport can be recovered from dynamic programming. For a terminal cost G , define

$$V(t, \mathbf{x}) = \inf_{\mathbf{u}} \left\{ \int_t^1 \frac{1}{2} \|\mathbf{u}(s)\|^2 ds + G(\mathbf{x}^{\mathbf{u}}(1)) \right\}.$$

Bellman's dynamic programming principle yields

$$0 = \inf_{\mathbf{u}} \left\{ \frac{1}{2} \|\mathbf{u}\|^2 + \partial_t V + \nabla V \cdot (\mathbf{A}(t)\mathbf{x} + \mathbf{B}(t)\mathbf{u}) \right\}.$$

Optimizing explicitly gives

$$\partial_t V + \nabla V \cdot \mathbf{A}(t)\mathbf{x} - \frac{1}{2} \nabla V^{\top} \mathbf{B}(t) \mathbf{B}(t)^{\top} \nabla V = 0, \quad (5.11)$$

with $V(1, \mathbf{y}) = G(\mathbf{y})$, and

$$\mathbf{u}^*(t, \mathbf{x}) = -\mathbf{B}(t)^{\top} \nabla V(t, \mathbf{x}).$$

An explicit potential is obtained from the **generalized Hopf–Lax formula**

$$\psi(t, \mathbf{x}) = \inf_{\mathbf{y}} \left\{ \psi(0, \mathbf{y}) + \frac{1}{2} (\mathbf{x} - \Phi(t, 0)\mathbf{y})^{\top} \mathbf{M}(t, 0)^{-1} (\mathbf{x} - \Phi(t, 0)\mathbf{y}) \right\},$$

with $V = -\psi$ and an initial potential determined by the Brenier potential of the transformed problem.

Schrödinger Bridges with a Markov Reference

Let $q(s, \mathbf{x}, t, \mathbf{y})$ be a positive Markov transition density. It satisfies the Chapman–Kolmogorov identity

$$q(t_1, \mathbf{x}_1, t_3, \mathbf{x}_3) = \int q(t_1, \mathbf{x}_1, t_2, \mathbf{x}_2) q(t_2, \mathbf{x}_2, t_3, \mathbf{x}_3) d\mathbf{x}_2. \quad (5.12)$$

Let $\mathbb{Q}$ be the corresponding Markov path law. Its endpoint distribution has density $q(0, \mathbf{x}, 1, \mathbf{y})$ relative to the chosen initial reference.

Theorem 5.3 The static Schrödinger bridge with positive Markov reference has endpoint law

$$d\mathbb{P}_{01}(\mathbf{x}, \mathbf{y}) = q(0, \mathbf{x}, 1, \mathbf{y}) \hat{\varphi}_0(\mathbf{x}) \varphi_1(\mathbf{y}) d\mathbf{x} d\mathbf{y} \quad (5.13)$$

for positive functions $\hat{\varphi}_0$ and φ_1 . The time- t marginal is

$$\rho_t(\mathbf{x}) = \varphi(t, \mathbf{x}) \hat{\varphi}(t, \mathbf{x}), \quad (5.14)$$

where

$$\varphi(t, \mathbf{x}) = \int q(t, \mathbf{x}, 1, \mathbf{y}) \varphi_1(\mathbf{y}) d\mathbf{y}$$

and

$$\hat{\varphi}(t, \mathbf{x}) = \int q(0, \mathbf{y}, t, \mathbf{x}) \hat{\varphi}_0(\mathbf{y}) d\mathbf{y}.$$

The two functions propagate backward and forward through the **Markov semigroup**. Equation (5.12) then turns the endpoint factorization (5.13) into the intermediate factorization (5.14). In a finite state space, this is the continuous time counterpart of the two Sinkhorn scaling vectors.

Return to the stochastic linear system (5.7). The **Eulerian control problem** is

$$\inf_{\rho, \mathbf{u}} \int_0^1 \int \frac{1}{2} \|\mathbf{u}(t, \mathbf{x})\|^2 \rho(t, \mathbf{x}) d\mathbf{x} dt$$

subject to the **Fokker–Planck equation**

$$\begin{aligned} \partial_t \rho + \nabla \cdot ((\mathbf{A}(t)\mathbf{x} + \mathbf{B}(t)\mathbf{u})\rho) - \frac{\varepsilon}{2} \sum_{i,j} \frac{\partial^2 (a_{ij}(t)\rho)}{\partial x_i \partial x_j} &= 0, \\ \rho(0, \cdot) &= \rho_0 \quad \text{and} \quad \rho(1, \cdot) = \rho_1, \end{aligned} \quad (5.15)$$

where $\mathbf{a}(t) = \mathbf{B}(t)\mathbf{B}(t)^\top$.

Theorem 5.4 Under suitable regularity conditions, the optimal control is

$$\mathbf{u}^*(t, \mathbf{x}) = \varepsilon \mathbf{B}(t)^\top \nabla \log \varphi(t, \mathbf{x}),$$

where φ is the backward factor from [Theorem 5.3](#) for the uncontrolled prior dynamics.

The drift is a **Doob h -transform**: the logarithmic gradient of the **backward likelihood** tilts the reference process just enough to attain the terminal marginal. When the endpoint laws are Gaussian and the prior is linear, every interpolating law remains Gaussian and the potentials are explicit. A second-order example,

$$d\mathbf{x}(t) = \mathbf{v}(t) dt \quad \text{and} \quad d\mathbf{v}(t) = \sqrt{\varepsilon} d\mathbf{w}(t) + \mathbf{u}(t) dt,$$

controls acceleration rather than position directly; the evolving distribution lives on **position–velocity phase space**.

Multimarginal, Unbalanced, and Martingale Transport

The **multimarginal extension** replaces a coupling of two measures by a joint law with prescribed marginals $\alpha_1, \dots, \alpha_S$:

$$\inf_{\gamma \in \Pi(\alpha_1, \dots, \alpha_S)} \int c(x_1, \dots, x_S) d\gamma. \quad (5.16)$$

An **entropic penalty** may be added exactly as in [\(3.3\)](#).

An economic example matches one consumer of type x_1 with contractors of types $x_2, \dots, x_S$. If a team produces project y , its negative welfare is

$$-p_1(x_1, y) + \sum_{i=2}^S c_i(x_i, y).$$

Optimizing the project within each team defines the **multimarginal cost**

$$c(x_1, \dots, x_S) = \inf_y \left\{ -p_1(x_1, y) + \sum_{i=2}^S c_i(x_i, y) \right\}.$$

Problem [\(5.16\)](#) then maximizes aggregate welfare by choosing the assignment of types.

Unbalanced transport relaxes the marginal constraints and penalizes their violation. For a non-

negative transport measure γ with marginals α and β , one minimizes

$$\inf_{\gamma \geq 0} \left\{ \int c \, d\gamma + \tau D_f(\alpha \mid \mu) + \tau D_f(\beta \mid \nu) \right\}. \quad (5.17)$$

This is valuable when total signal strength matters, normalization is not meaningful, or the observed marginals contain noise. In such settings, enforcing exact mass conservation can amplify error.

Martingale transport adds a conditional expectation constraint. Let $\mu_1, \dots, \mu_N$ be probability measures on $\mathbb{R}$ in **increasing convex order**, meaning that

$$t \mapsto \int \phi \, d\mu_t$$

is nondecreasing for every convex ϕ . This is necessary and sufficient for them to be the marginals of a martingale (X_t) satisfying

$$\mathbb{E}[X_t \mid X_1, \dots, X_s] = X_s, \quad s \leq t.$$

The **martingale optimal transport problem** minimizes $\int c \, d\gamma$ over all such joint laws. In finance, c may be the payoff of a **path-dependent option**. The infimum and supremum then give bounds on arbitrage-free prices without specifying a model, while the dual variables describe hedging strategies.

Transport can also be defined for positive matrices, possibly of unit trace. Unlike a scalar histogram, a matrix records both eigenvalues and eigenvectors. Relevant costs include the Bures–Wasserstein metric and **matrix relative entropy**

$$\text{KL}(\mathbf{A} \mid \mathbf{B}) = \text{tr}(\mathbf{A} \log \mathbf{A} - \mathbf{A} \log \mathbf{B} - \mathbf{A} + \mathbf{B}).$$

This points toward quantum and matrix analogues of the Dirichlet transport and $L^{(\alpha)}$ -divergence.

Gromov–Wasserstein Geometry

When two measures live on different spaces, a pointwise cost between their locations may not be defined. **Gromov-type distances** compare the internal geometries instead.

For nonempty compact subsets A, B of a metric space (Z, d_Z) , the **Hausdorff distance** is

$$H_Z(A, B) = \max \left\{ \sup_{a \in A} \inf_{b \in B} d_Z(a, b), \sup_{b \in B} \inf_{a \in A} d_Z(a, b) \right\}.$$

The **Gromov–Hausdorff distance** between compact metric spaces (X, d_X) and (Y, d_Y) is

$$GH(X, Y) = \inf_{Z, f, g} H_Z(f(X), g(Y)), \quad (5.18)$$

where $f: X \rightarrow Z$ and $g: Y \rightarrow Z$ are isometric embeddings.

An equivalent formulation uses a correspondence.

Definition 5.5 A **correspondence** between X and Y is a subset $R \subseteq X \times Y$ whose projections onto both factors are surjective.

Then

$$GH(X, Y) = \frac{1}{2} \inf_R \sup_{(x, y), (x', y') \in R} |d_X(x, x') - d_Y(y, y')|. \quad (5.19)$$

The supremum is the **distortion** of the correspondence.

Figure 12 compares two configurations by their internal distances alone.

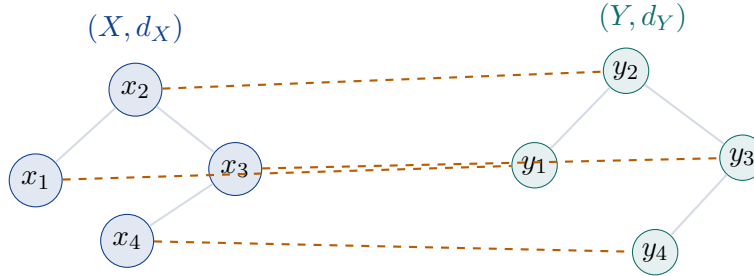


FIGURE 12

For metric-measure spaces (X, d_X, μ_X) and (Y, d_Y, μ_Y) , replace a set correspondence by a probabilistic coupling.

Definition 5.6 The squared **Gromov–Wasserstein distance** is

$$GW(X, Y)^2 = \inf_{\gamma \in \Pi(\mu_X, \mu_Y)} \iint |d_X(x, x') - d_Y(y, y')|^2 d\gamma(x, y) d\gamma(x', y'). \quad (5.20)$$

This distance is invariant under measure-preserving isometries. It is a smooth probabilistic analogue of (5.19) and is useful in shape matching, where objects may have no shared coordinate system.

Lecture 6 — Friday, October 18, 2019

5. Variational Problems and Statistical Learning

Barycentres and Exponential Families

Given $\mathbf{x}_1, \dots, \mathbf{x}_N \in \mathbb{R}^d$ and weights $\lambda_i \geq 0$ with $\sum_i \lambda_i = 1$, the **weighted Euclidean mean**

$$\bar{\mathbf{x}}_\lambda = \sum_{i=1}^N \lambda_i \mathbf{x}_i$$

is the unique minimizer of

$$\mathbf{x} \mapsto \sum_{i=1}^N \lambda_i \|\mathbf{x} - \mathbf{x}_i\|^2.$$

A **barycentre** replaces the squared Euclidean distance by another divergence or metric.

Theorem 6.1 Let D_ϕ be the Bregman divergence from [Definition 2.10](#). Its **right barycentre** is the arithmetic mean:

$$\operatorname{argmin}_{\mathbf{s}} \sum_{i=1}^N \lambda_i D_\phi[\mathbf{x}_i : \mathbf{s}] = \sum_{i=1}^N \lambda_i \mathbf{x}_i.$$

Proof. Set $\bar{\mathbf{x}}_\lambda = \sum_i \lambda_i \mathbf{x}_i$. Expanding the definition of the Bregman divergence gives

$$\begin{aligned} \sum_{i=1}^N \lambda_i D_\phi[\mathbf{x}_i : \mathbf{s}] &= \sum_{i=1}^N \lambda_i (\phi(\mathbf{x}_i) - \phi(\mathbf{s}) - \nabla \phi(\mathbf{s}) \cdot (\mathbf{x}_i - \mathbf{s})) \\ &= \sum_{i=1}^N \lambda_i (\phi(\mathbf{x}_i) - \phi(\bar{\mathbf{x}}_\lambda) + D_\phi[\bar{\mathbf{x}}_\lambda : \mathbf{s}]). \end{aligned}$$

The first two terms do not depend on $\mathbf{s}$. The final term is nonnegative and, by strict convexity, vanishes only when $\mathbf{s} = \bar{\mathbf{x}}_\lambda$. Thus the weighted mean is the unique minimizer. $\square$

This observation underlies **Bregman clustering**; see [Banerjee, Merugu, Dhillon, and Ghosh \(2005\)](#).

Now consider an exponential family

$$p(x; \boldsymbol{\theta}) = \exp(\boldsymbol{\theta} \cdot \mathbf{h}(x) - \phi(\boldsymbol{\theta})) \quad (6.1)$$

relative to a fixed base measure. The log-partition function ϕ is convex, and

$$D_\phi[\boldsymbol{\theta} : \boldsymbol{\theta}'] = \text{KL}(p(\cdot; \boldsymbol{\theta}') \mid p(\cdot; \boldsymbol{\theta})).$$

Its convex conjugate is **negative Shannon entropy**,

$$\psi(\boldsymbol{\eta}) = -H(p(\cdot; \boldsymbol{\theta})) = \int p(x; \boldsymbol{\theta}) \log p(x; \boldsymbol{\theta}) \, d\mu(x),$$

where

$$\boldsymbol{\eta} = \nabla \phi(\boldsymbol{\theta}) = \mathbb{E}_{\boldsymbol{\theta}}[\mathbf{h}(X)]$$

is the expectation coordinate. If $\boldsymbol{\eta}' = \nabla \phi(\boldsymbol{\theta}')$, then

$$D_\phi[\boldsymbol{\theta} : \boldsymbol{\theta}'] = D_\psi[\boldsymbol{\eta}' : \boldsymbol{\eta}] = \phi(\boldsymbol{\theta}) + \psi(\boldsymbol{\eta}') - \boldsymbol{\theta} \cdot \boldsymbol{\eta}'. \quad (6.2)$$

For a data point x for which $\mathbf{h}(x)$ lies in the effective domain of ψ , take $\boldsymbol{\eta}' = \mathbf{h}(x)$. Equations (6.1) and (6.2) give

$$\log p(x; \boldsymbol{\theta}) = -D_\psi[\mathbf{h}(x) : \boldsymbol{\eta}] + \psi(\mathbf{h}(x)).$$

Therefore maximum likelihood for observations $x_1, \dots, x_N$ is equivalent to

$$\underset{\boldsymbol{\eta}}{\operatorname{argmin}} \sum_{i=1}^N D_\psi[\mathbf{h}(x_i) : \boldsymbol{\eta}],$$

and [Theorem 6.1](#) gives

$$\hat{\boldsymbol{\eta}} = \frac{1}{N} \sum_{i=1}^N \mathbf{h}(x_i).$$

Thus the maximum likelihood estimator is a right Bregman barycentre in expectation coordinates.

Wasserstein Barycentres

Definition 6.2 Let $\nu_1, \dots, \nu_N \in \mathcal{P}_2(\mathbb{R}^d)$ and let $\lambda_i \geq 0$ with $\sum_i \lambda_i = 1$. A **Wasserstein barycentre** is a minimizer of

$$\inf_{\nu \in \mathcal{P}_2(\mathbb{R}^d)} \sum_{i=1}^N \lambda_i W_2(\nu_i, \nu)^2. \quad (6.3)$$

Theorem 6.3 A Wasserstein barycentre exists. It is unique if at least one positive-weight measure ν_i vanishes on every Borel set of Hausdorff dimension at most $d - 1$; in particular, absolute continuity of one positive-weight ν_i is sufficient.

This result is due to [Agueh and Carlier \(2011\)](#).

Zero-weight inputs may be discarded, so assume in the following dual formulas that every $\lambda_i > 0$. The dual problem uses a **weighted quadratic infimal convolution**. Define

$$S_\lambda f(\mathbf{x}) = \inf_{\mathbf{y} \in \mathbb{R}^d} \left\{ \frac{\lambda}{2} \|\mathbf{x} - \mathbf{y}\|^2 - f(\mathbf{y}) \right\}.$$

With a harmless factor of $1/2$ in the primal objective, the dual is

$$\sup_{\substack{f_1, \dots, f_N \\ \sum_i f_i = 0}} \sum_{i=1}^N \int S_{\lambda_i} f_i \, d\nu_i. \quad (6.4)$$

Indeed,

$$S_{\lambda_i} f_i(\mathbf{x}) + f_i(\mathbf{y}) \leq \frac{\lambda_i}{2} \|\mathbf{x} - \mathbf{y}\|^2.$$

Integrating this inequality against couplings of ν_i and a candidate barycentre ν , then summing over i , eliminates the functions f_i because $\sum_i f_i = 0$. Convex duality turns the resulting weak inequality into equality.

There is also a **multimarginal formulation**. For $\mathbf{x} = (\mathbf{x}_1, \dots, \mathbf{x}_N)$, set

$$T(\mathbf{x}) = \sum_{i=1}^N \lambda_i \mathbf{x}_i.$$

Consider

$$\inf_{\gamma \in \Pi(\nu_1, \dots, \nu_N)} \int \sum_{i=1}^N \frac{\lambda_i}{2} \|\mathbf{x}_i - T(\mathbf{x})\|^2 d\gamma(\mathbf{x}). \quad (6.5)$$

Theorem 6.4 If γ solves (6.5), then

$$\nu = T_{\#}\gamma$$

solves (6.3).

Proof. Let $\pi_i(\mathbf{x}) = \mathbf{x}_i$. Then

$$\gamma_i = (\pi_i, T)_{\#}\gamma \in \Pi(\nu_i, \nu),$$

and hence

$$W_2(\nu_i, \nu)^2 \leq \int \|\mathbf{x}_i - T(\mathbf{x})\|^2 d\gamma(\mathbf{x}).$$

Summing with weights proves one inequality between the two problems. Conversely, glue optimal couplings of a candidate ν with each ν_i over their common ν -coordinate, just as in the proof of Theorem 2.5. Conditional on a barycentre point $\mathbf{y}$, the weighted mean minimizes

$$\mathbf{z} \mapsto \sum_i \lambda_i \|\mathbf{x}_i - \mathbf{z}\|^2.$$

This gives the reverse inequality and shows that a multimarginal optimizer pushes forward to a barycentre. $\square$

In one dimension, **quantiles** give an explicit solution. If F_i and F are the **distribution functions** of ν_i and ν , then

$$W_2(\nu_i, \nu)^2 = \int_0^1 |F_i^{-1}(u) - F^{-1}(u)|^2 du.$$

Pointwise minimization yields

$$F^{-1}(u) = \sum_{i=1}^N \lambda_i F_i^{-1}(u). \quad (6.6)$$

If ν_1 is atomless, this common-quantile coupling can equivalently be written as $\nu = T_{\#}\nu_1$, with

$$T = \sum_{i=1}^N \lambda_i T_{1i} \quad \text{and} \quad (T_{1i})_{\#}\nu_1 = \nu_i.$$

For equally weighted empirical measures with the same number of atoms, (6.6) averages atoms of **equal rank**, as illustrated in Figure 13.

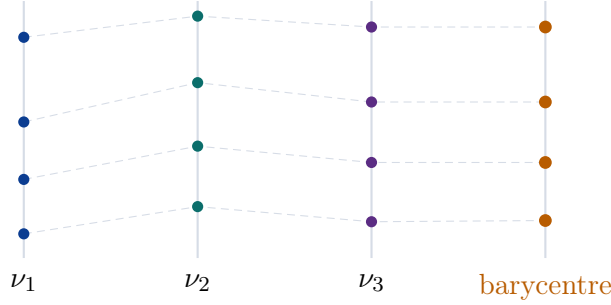


FIGURE 13

When $N = 2$, write the inputs as ν_0 and ν_1 . If ν_0 is absolutely continuous, the barycentre with weights $(1 - t, t)$ is precisely the displacement interpolation

$$\nu_t = ((1 - t) \text{Id} + tT)_{\#} \nu_0.$$

For centered Gaussian inputs with positive-definite covariance matrices $\nu_i = \mathcal{N}(\mathbf{0}, \Sigma_i)$, the barycentre is Gaussian, $\nu = \mathcal{N}(\mathbf{0}, \Sigma)$, where

$$\sum_{i=1}^N \lambda_i (\Sigma^{1/2} \Sigma_i \Sigma^{1/2})^{1/2} = \Sigma. \quad (6.7)$$

This matrix equation is the covariance analogue of averaging aligned quantiles.

In the discrete setting, an **entropic barycentre** solves

$$\min_{\mathbf{a}} \sum_{i=1}^N \lambda_i L^{\varepsilon}(\mathbf{a}, \mathbf{b}_i). \quad (6.8)$$

For $\varepsilon > 0$, L^{ε} is convex and differentiable on the interior of the product of simplices. Its gradient with respect to the marginals is given by the **optimal dual potentials** $(\mathbf{f}, \mathbf{g})$, where

$$L^{\varepsilon}(\mathbf{a}, \mathbf{b}) = \max_{\mathbf{f}, \mathbf{g}} \{ \langle \mathbf{f}, \mathbf{a} \rangle + \langle \mathbf{g}, \mathbf{b} \rangle - \varepsilon \langle \exp(\mathbf{f}/\varepsilon), \mathbf{K} \exp(\mathbf{g}/\varepsilon) \rangle + \varepsilon \}.$$

The **Sinkhorn scalings** satisfy

$$\mathbf{u} = \exp(\mathbf{f}/\varepsilon) \quad \text{and} \quad \mathbf{v} = \exp(\mathbf{g}/\varepsilon).$$

At a solution of (6.8), the **first order condition** is $\sum_i \lambda_i \mathbf{f}_i = \mathbf{0}$ up to the common **additive gauge**. This leads to both gradient methods and **generalized Sinkhorn iterations**. In imaging applications, the barycentre blends corresponding shapes and colours rather than averaging pixels at fixed locations.

The population version also has a **law of large numbers**. Let $\tilde{\mathbb{P}}$ be a probability distribution on $\mathcal{P}_2(X)$ and suppose the **population barycentre**

$$\nu^* = \operatorname{argmin}_{\nu} \int W_2(\mu, \nu)^2 d\tilde{\mathbb{P}}(\mu)$$

is unique.

Theorem 6.5 If $\mu_1, \dots, \mu_N$ are independent samples from $\tilde{\mathbb{P}}$ and ν_N is their **empirical Wasserstein barycentre**, then, under the required moment and compactness assumptions,

$$W_2(\nu_N, \nu^*) \longrightarrow 0$$

almost surely as $N \rightarrow \infty$.

Metric Gradient Flows

For a differentiable function F on Euclidean space, the **gradient flow**

$$\dot{\mathbf{x}}(t) = -\nabla F(\mathbf{x}(t))$$

is the continuous limit of the **proximal recursion**

$$\mathbf{x}^{(\ell+1)} = \operatorname{argmin}_{\mathbf{x}} \left\{ \frac{1}{2} \|\mathbf{x} - \mathbf{x}^{(\ell)}\|^2 + \tau F(\mathbf{x}) \right\}. \quad (6.9)$$

On a Riemannian manifold, the squared norm is replaced by squared Riemannian distance. This observation makes sense even in a metric space (X, d) : define a candidate flow as the small- τ limit of

$$x^{(\ell+1)} = \operatorname{argmin}_{x \in X} \left\{ \frac{1}{2} d(x, x^{(\ell)})^2 + \tau F(x) \right\}. \quad (6.10)$$

This is the **minimizing movement construction**.

The heat equation gives two complementary examples. First consider the **Dirichlet energy**

$$F(u) = \frac{1}{2} \int_{\mathbb{R}^n} \|\nabla u\|^2 d\mathbf{x}.$$

For a smooth compactly supported variation v ,

$$\begin{aligned} F(u + hv) &= \frac{1}{2} \int \|\nabla u + h\nabla v\|^2 d\mathbf{x} \\ &= F(u) + h \int \nabla u \cdot \nabla v d\mathbf{x} + O(h^2) \\ &= F(u) - h \int (\Delta u)v d\mathbf{x} + O(h^2). \end{aligned}$$

With the L^2 inner product, $\text{grad } F(u) = -\Delta u$. Hence

$$\partial_t u = \Delta u$$

is the **L^2 gradient flow** of Dirichlet energy.

The same heat equation is the **Wasserstein gradient flow** of entropy. Let F be a smooth functional of densities, and denote its **first variation** by $\delta F / \delta \rho$. A tangent in Otto's metric has the form

$$\sigma = -\nabla \cdot (\rho \nabla \varphi).$$

Along a curve with this tangent,

$$\left. \frac{d}{dt} F(\rho_t) \right|_{t=0} = \int \frac{\delta F}{\delta \rho} \sigma d\mathbf{x} = \int \nabla \cdot \left(\frac{\delta F}{\delta \rho} \right) \cdot \nabla \varphi d\rho.$$

Comparing with (4.1) identifies

$$\text{grad}_W F(\rho) = -\nabla \cdot \left(\rho \nabla \frac{\delta F}{\delta \rho} \right). \quad (6.11)$$

Therefore the negative gradient flow is

$$\partial_t \rho = \nabla \cdot \left(\rho \nabla \frac{\delta F}{\delta \rho} \right). \quad (6.12)$$

For negative entropy,

$$F(\rho) = \int \rho \log \rho d\mathbf{x},$$

the first variation is $\log \rho + 1$. The additive constant disappears after taking a spatial gradient, so

$$\nabla \cdot \left(\rho \nabla \frac{\delta F}{\delta \rho} \right) = \nabla \cdot (\rho \nabla \log \rho) = \Delta \rho.$$

Equation (6.12) is again the heat equation.

More generally, the Fokker–Planck equation

$$\partial_t \rho = \nabla \cdot (\rho \nabla \Psi) + \beta^{-1} \Delta \rho \tag{6.13}$$

is the Wasserstein gradient flow of the **free energy**

$$F(\rho) = \int \Psi \rho \, d\mathbf{x} + \beta^{-1} \int \rho \log \rho \, d\mathbf{x}.$$

It is also the forward equation of the **Langevin diffusion**

$$d\mathbf{X}_t = -\nabla \Psi(\mathbf{X}_t) \, dt + \sqrt{2\beta^{-1}} \, d\mathbf{W}_t.$$

When the normalizing integral is finite, the **stationary density** is proportional to $e^{-\beta \Psi}$.

For $m > 1$, another example is the **porous medium equation with drift**,

$$\partial_t \rho - \Delta \rho^m - \nabla \cdot (\rho \nabla V) = 0,$$

whose energy is

$$F(\rho) = \frac{1}{m-1} \int \rho^m \, d\mathbf{x} + \int V(\mathbf{x}) \rho(\mathbf{x}) \, d\mathbf{x}.$$

The rigorous theory is developed by [Ambrosio, Gigli, and Savaré \(2008\)](#). The entropy interpretation of the Fokker–Planck equation originates with [Jordan, Kinderlehrer, and Otto \(1998\)](#).

Particle Discretizations

The minimizing movement scheme (6.10) for the Wasserstein metric becomes

$$\rho^{(\ell+1)} = \operatorname{argmin}_{\rho} \left\{ \frac{1}{2} W_2(\rho, \rho^{(\ell)})^2 + \tau F(\rho) \right\}.$$

One numerical approximation represents the density by a **moving cloud**

$$\rho_t \approx \frac{1}{n} \sum_{i=1}^n \delta_{\mathbf{x}_i(t)}.$$

For sufficiently small time steps, the identity matching remains optimal, and

$$W_2(\rho_{t+\Delta t}, \rho_t)^2 = \frac{1}{n} \sum_{i=1}^n \|\mathbf{x}_i(t + \Delta t) - \mathbf{x}_i(t)\|^2.$$

Thus the measure-valued flow becomes an ordinary Euclidean gradient flow for the particle vector

$$\mathbf{X}(t) = (\mathbf{x}_1(t), \dots, \mathbf{x}_n(t)) \quad \text{and} \quad \dot{\mathbf{X}}(t) = -\nabla \tilde{F}(\mathbf{X}(t)).$$

For the Langevin example, the particles initially spread and then settle into the **Gibbs density**. The approximation converts diffusion into a deterministic interaction among particle locations through the **discretized entropy**.

The energy landscape in Figure 14 shows how each proximal step balances a short Wasserstein move against a decrease in free energy.

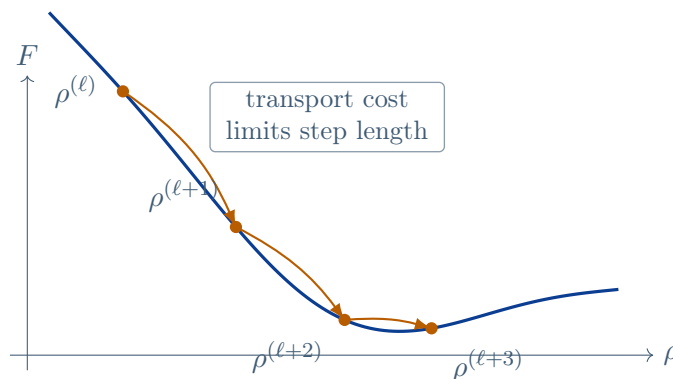


FIGURE 14

Minimum Kantorovich Estimation and Generative Models

Let $x_1, \dots, x_n$ be independent observations from an unknown distribution, and let

$$\beta_n = \frac{1}{n} \sum_{i=1}^n \delta_{x_i}$$

be the empirical measure. A **parametric model** is a map $\boldsymbol{\theta} \mapsto \alpha_{\boldsymbol{\theta}} \in \mathcal{P}(X)$. A general **minimum-distance estimator** solves

$$\min_{\boldsymbol{\theta} \in \Theta} \mathcal{L}(\alpha_{\boldsymbol{\theta}}, \beta_n). \tag{6.14}$$

If α_θ has density p_θ and

$$\mathcal{L}(\alpha_\theta, \beta_n) = - \int \log p_\theta(x) d\beta_n(x) = -\frac{1}{n} \sum_{i=1}^n \log p_\theta(x_i),$$

then (6.14) is maximum likelihood estimation. At the population level, if the true law β and the model law have densities p and p_θ with respect to a common dominating measure, then

$$\text{KL}(\beta \mid \alpha_\theta) = \int p(x) \log \left(\frac{p(x)}{p_\theta(x)} \right) dx.$$

The maximum likelihood loss differs from this relative entropy by a term independent of θ .

Relative entropy becomes infinite unless the data law is absolutely continuous with respect to the model law. This is a serious problem when both distributions are concentrated on low dimensional sets. In high dimensional data analysis, a distribution may lie on an unknown manifold M with $\dim M \ll d$; two nearby model manifolds can then have disjoint supports even when they are geometrically close. Moreover, the model density may be computationally inaccessible.

A generative model avoids explicit densities. Let ζ be a simple distribution on a low dimensional latent space Z , and let $h_\theta: Z \rightarrow X$ be a parameterized map (often a neural network). The model is

$$\alpha_\theta = (h_\theta)_\# \zeta. \quad (6.15)$$

Samples are easy to generate even if a density for (6.15) is unavailable. A **weak transport discrepancy** is therefore a natural loss.

The standard comparison uses **total variation**, relative entropy, **Jensen–Shannon divergence**, and Wasserstein distance:

$$\begin{aligned} \delta(\mathbb{P}_r, \mathbb{P}_g) &= \sup_A |\mathbb{P}_r(A) - \mathbb{P}_g(A)|, \\ \text{KL}(\mathbb{P}_r \mid \mathbb{P}_g) &= \int \log \left(\frac{d\mathbb{P}_r}{d\mathbb{P}_g} \right) d\mathbb{P}_r, \\ JS(\mathbb{P}_r, \mathbb{P}_g) &= \frac{1}{2} \text{KL}(\mathbb{P}_r \mid \mathbb{P}_m) + \frac{1}{2} \text{KL}(\mathbb{P}_g \mid \mathbb{P}_m) \quad \text{and} \quad \mathbb{P}_m = \frac{1}{2}(\mathbb{P}_r + \mathbb{P}_g). \end{aligned}$$

Example 6.6 Let Z be uniform on $[0, 1]$. Let $\mathbb{P}_0$ be the law of $(0, Z) \in \mathbb{R}^2$, and let $\mathbb{P}_\theta$ be the law of (θ, Z) . Then

$$W_1(\mathbb{P}_0, \mathbb{P}_\theta) = |\theta|,$$

whereas

$$JS(\mathbb{P}_0, \mathbb{P}_\theta) = \begin{cases} \log 2, & \theta \neq 0, \\ 0, & \theta = 0, \end{cases} \quad \text{and} \quad KL(\mathbb{P}_0 \mid \mathbb{P}_\theta) = \begin{cases} \infty, & \theta \neq 0, \\ 0, & \theta = 0, \end{cases}$$

and

$$\delta(\mathbb{P}_0, \mathbb{P}_\theta) = \begin{cases} 1, & \theta \neq 0, \\ 0, & \theta = 0. \end{cases}$$

The geometry in Figure 15 is simple: Wasserstein distance records the horizontal displacement, while divergences that are sensitive to the support jump as soon as the lines separate.

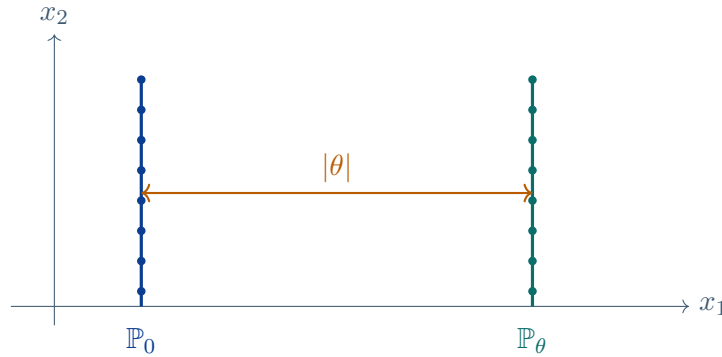


FIGURE 15

Theorem 6.7 Let $\mathbb{P}_r$ be fixed, let Z have a fixed latent distribution, and let $g: Z \times \mathbb{R}^d \rightarrow X$ define $\mathbb{P}_\theta$ as the law of $g(Z, \theta)$.

1. If g is continuous in θ and the required moments are locally dominated, then $W_p(\mathbb{P}_r, \mathbb{P}_\theta)$ is continuous in θ .
2. If g is locally Lipschitz and satisfies the corresponding regularity condition, then $W_p(\mathbb{P}_r, \mathbb{P}_\theta)$ is continuous everywhere and differentiable almost everywhere.
3. The analogous assertions fail in general for Jensen–Shannon divergence and both orientations of relative entropy.

The theorem, together with Example 6.6, explains why **Wasserstein loss** can support **training by gradient methods** even when model and data distributions occupy different low dimensional sets.

For the cost $c(x, y) = d(x, y)$, c -concave potentials are exactly 1-Lipschitz functions, and $f^c = -f$. [Kantorovich duality](#) therefore becomes the [Kantorovich–Rubinstein formula](#)

$$W_1(\mu, \nu) = \sup_{\|f\|_{\text{Lip}} \leq 1} \left\{ \int f \, d\mu - \int f \, d\nu \right\}. \quad (6.16)$$

For a generative model $\mathbb{P}_\theta = (g_\theta)_\# \zeta$,

$$W_1(\mathbb{P}_r, \mathbb{P}_\theta) = \sup_{\|f\|_{\text{Lip}} \leq 1} \{ \mathbb{E}_{\mathbb{P}_r}[f(X)] - \mathbb{E}_\zeta[f(g_\theta(Z))] \}.$$

The original **Wasserstein generative adversarial network** parameterizes f by a **critic network** f_w and solves the saddle problem

$$\min_{\theta} \max_w \{ \mathbb{E}_{\mathbb{P}_r}[f_w(X)] - \mathbb{E}_\zeta[f_w(g_\theta(Z))] \}, \quad (6.17)$$

subject to an approximate Lipschitz constraint on the critic. This formulation was introduced by [Arjovsky, Chintala, and Bottou \(2017\)](#).

Appendix: Lecture and Tutorial Index

1. Lecture 1 — Friday, September 13, 2019
2. Lecture 2 — Friday, September 20, 2019
3. Lecture 3 — Friday, September 27, 2019
4. Lecture 4 — Friday, October 04, 2019
5. Lecture 5 — Friday, October 11, 2019
6. Lecture 6 — Friday, October 18, 2019