

# Generalized Gradient Flows for Jump Processes

Oliver Tse

Eindhoven University of Technology

Updated on September 4, 2026

Lecture notes for a school on Gradient Flows at ICTS

*Abstract.* These notes accompany a short course on *generalized gradient structures* for Markov jump processes. Starting from reversible Markov chains on discrete state spaces, we see that one equation may carry many gradient structures, or none, and we make the *cosh*-type structure rigorous for jump processes on Polish spaces. Which structure to take is settled by a selection principle: the large-deviation rate functional of the underlying particle system is canonical, and splitting it into its symmetric and antisymmetric parts under time reversal produces a dissipation and a free energy—the reason jump processes come out *cosh*.

## Contents

|          |                                                                                            |           |
|----------|--------------------------------------------------------------------------------------------|-----------|
| <b>1</b> | <b>Gradient structures for discrete Markov chains</b>                                      | <b>1</b>  |
| 1.1      | Warm-up: gradient flows on Euclidean space . . . . .                                       | 1         |
| 1.2      | Beyond quadratic dissipation potentials . . . . .                                          | 3         |
| 1.3      | Reversible Markov chains as gradient flows . . . . .                                       | 4         |
| 1.4      | The Fisher information and dissipation along solutions . . . . .                           | 7         |
| <b>2</b> | <b>Generalized gradient structures for jump processes</b>                                  | <b>8</b>  |
| 2.1      | The three ingredients of the gradient structure . . . . .                                  | 9         |
| 2.2      | The Fisher information . . . . .                                                           | 11        |
| 2.3      | Admissible curves and the chain rule . . . . .                                             | 12        |
| 2.4      | Definition of solutions: the $(\mathcal{F}, \mathcal{R}, \mathcal{R}^*)$ balance . . . . . | 13        |
| 2.5      | Existence via minimizing movements . . . . .                                               | 13        |
| <b>3</b> | <b>Where gradient structures come from</b>                                                 | <b>14</b> |
| 3.1      | The selection principle: large deviations . . . . .                                        | 14        |
| 3.2      | Time reversal performs the splitting . . . . .                                             | 17        |

# 1 Gradient structures for discrete Markov chains

ch:discrete

## Goal of Lecture 1

Starting from the classical Euclidean picture, we explain what a (*generalized*) *gradient structure* is, and we see that one equation may carry many such structures, or none at all. Onsager’s reciprocal relations explain why any of them is a convex duality between conjugate dissipation potentials—though not which pair to take. We then show that a reversible continuous-time Markov chain on a finite state space is a gradient flow of the relative entropy. We meet the discrete continuity equation, the *graph gradient/graph divergence* duality, and the freedom in choosing a *dissipation potential*—in particular the quadratic (Maas–Mielke–Chow–Huang–Zhou) and the *cosh* structures.

## 1.1 Warm-up: gradient flows on Euclidean space

sec:euclidean

Consider the ODE

$$\dot{x}(t) = -\nabla F(x(t)), \quad x(0) = x_0 \in \mathbb{R}^d, \quad (1.1)$$

eq:gf-euclid

with  $F: \mathbb{R}^d \rightarrow \mathbb{R}$  smooth. Calling (1.1) a *gradient flow* is, strictly speaking, an incomplete statement: it becomes meaningful only once we fix a *gradient structure*, consisting of

- (i) a state space  $\mathcal{X}$  (here  $\mathbb{R}^d$ );
- (ii) a driving energy  $F: \mathcal{X} \rightarrow \mathbb{R} \cup \{+\infty\}$ ;
- (iii) notions of velocity  $\dot{x}$  and of gradient  $\nabla F$ ;
- (iv) a *kinetic relation* linking the restoring force  $-\nabla F(x)$  to the velocity  $\dot{x}$ .

The same equation can carry several, or even infinitely many, gradient structures, each with possibly different physical or probabilistic meaning.

**The kinetic relation.** For the moment, we consider (iv) in the linear form

$$\dot{x}(t) = -K(x(t)) \nabla F(x(t)), \quad (GF)$$

eq:kinetic-linear

where the *mobility*, or *Onsager operator*  $K(x)$  is symmetric and positive definite. Its inverse  $G(x) := K(x)^{-1}$  is then a metric tensor, and  $K = \text{Id}$  returns (1.1). The pair carries with it two quadratic forms, the *dissipation potential* and its *dual*:

$$\mathcal{R}(x, v) := \frac{1}{2} \langle v, G(x)v \rangle = \frac{1}{2} |v|_G^2, \quad \mathcal{R}^*(x, \xi) := \frac{1}{2} \langle \xi, K(x)\xi \rangle = \frac{1}{2} |\xi|_K^2,$$

where the first measures a velocity  $v$ , while the second a force  $\xi$ .

**Remark 1.1** (Onsager reciprocity) The symmetry and positive definiteness of  $G$  are not conveniences but physics. They are Onsager’s *reciprocal relations* from “Reciprocal relations in irreversible processes” (1931), which earned him the 1963 Nobel Prize. In modern language, they are precisely the statement  $G^\top = G$ . The symmetry of the Onsager operator is the macroscopic shadow of *microscopic reversibility* of the underlying dynamics—a first hint that reversibility and gradient structure belong together, which Lecture 3 will make precise.

**Variational reformulation.** The defining characteristic that makes a gradient flow a gradient flow is that it dissipates a driving energy along the flow, i.e., along a solution of (GF) the *chain rule* gives

$$\frac{d}{dt}F(x) = \langle \nabla F(x), \dot{x} \rangle = -\langle \nabla F(x), K(x)\nabla F(x) \rangle \leq 0.$$

Yet this identity provides a stronger characterization of gradient flows. Indeed, from the Cauchy–Schwarz and Young inequalities, we obtain

$$\langle \xi, v \rangle \leq \|v\|_G \|\xi\|_K \leq \frac{1}{2} \|v\|_G^2 + \frac{1}{2} \|\xi\|_K^2 = \mathcal{R}(x, v) + \mathcal{R}^*(x, \xi)$$

for any velocity  $v$  and force  $\xi$ , with equality if and only if  $v = K(x)\xi$ . Therefore, taking  $\xi = -\nabla F(x)$  and  $v = \dot{x}$ , this allows us to fully characterize the *vector* equation (GF) with a single *scalar* one:

$$\mathcal{R}(x, \dot{x}) + \mathcal{R}^*(x, -\nabla F(x)) = \langle \dot{x}, -\nabla F(x) \rangle = -\frac{d}{dt}F(x(t)). \quad \text{eq:onsager-edp} \quad (1.2)$$

Integrating over  $[s, t]$  yields the so-called *energy–dissipation balance* (EDB)

$$\mathcal{L}(x, \dot{x}; [s, t]) := \int_s^t \left( \mathcal{R}(x, \dot{x}) + \mathcal{R}^*(x, -\nabla F(x)) \right) dr + F(x(t)) - F(x(s)) = 0. \quad \text{eq:edh-euclid} \quad (\text{EDB})$$

Since the functional  $\mathcal{L}$  is always nonnegative due to the chain rule and Young’s inequality, the equality  $\mathcal{L} = 0$  on every interval  $[s, t]$  characterizes solutions of (GF). With  $K = \text{Id}$  the two potentials are the familiar  $\frac{1}{2} |\dot{x}|^2$  and  $\frac{1}{2} |\nabla F(x)|^2$ .

Observe that (EDB) is purely *variational* and no longer refers to a Hilbert structure: it needs only a notion of *metric speed*  $|\dot{x}|$  and of *metric slope*  $|\nabla F|$ , and therefore extends to general metric spaces and to spaces of measures.

**Some examples.** Let us provide concrete illustrations. Recall that a gradient structure consists of a state space, a driving energy, notions of velocity and gradient, and a kinetic relation, but the equation alone determines none of these objects. The point of the example below is that the vector field alone does not determine the gradient structure.

**Example 1.2** (One equation, two structures) On  $\mathcal{X} = \mathbb{R}^2$  consider

$$\dot{x}(t) = V(x(t)), \quad V(x) = -(x_1, x_2 + \alpha x_2^3)^\top, \quad \alpha > 0.$$

Setting  $F(x) = \frac{1}{2}(x_1^2 + x_2^2) + \frac{\alpha}{4}x_2^4$  gives  $\nabla F = -V$ , so this solves (GF) with

$$F(x) = \frac{1}{2}(x_1^2 + x_2^2) + \frac{\alpha}{4}x_2^4, \quad K(x) = \text{Id}.$$

But it is *equally* a gradient flow with

$$\tilde{F}(x) = \frac{1}{2}(x_1^2 + x_2^2), \quad \tilde{K}(x) = \begin{pmatrix} 1 & 0 \\ 0 & 1 + \alpha x_2^2 \end{pmatrix},$$

since  $\tilde{K}(x)\nabla\tilde{F}(x) = (x_1, x_2 + \alpha x_2^3)^\top = -V(x)$ .

Therefore, just looking at the ODE, one cannot tell whether the nonlinearity  $\alpha x_2^3$  comes from a non-quadratic *energy* or from a state-dependent *friction law*. The choice of a gradient structure is a modeling choice, and it carries strictly more information than the equation itself.

ex:damped

**Example 1.3** (An equation with no gradient structure) The damped harmonic oscillator

$$\dot{x}(t) = V(x(t)), \quad V(x) = (x_2, -x_1 - \gamma x_2)^\top, \quad \gamma > 0,$$

possesses no gradient structure at all. It does carry a more general structure—the so-called GENERIC structure, combining a conservative (Hamiltonian) part with a dissipative one—which remains an active field of research. Dissipating energy is thus necessary but not sufficient: the oscillator’s energy decays, yet the motion is not steepest descent for any choice of friction.

The moral of the story here is that an evolution  $\dot{x} = V(x)$  may carry many gradient structures, or none. Either way, which one to take, if a gradient structure exists, is a question the equation alone cannot answer. We will discuss this in detail in Lecture 3.

## 1.2 Beyond quadratic dissipation potentials

sec:generalized

So far, we’ve only considered quadratic dissipation potentials. However, nothing in EDB actually requires this. Relaxing it is exactly what jump dynamics will force us to do in §3. To exploit that freedom, and to reach the (generalized) gradient structures for jump dynamics, we will need to replace the quadratic Young inequality by its convex-duality generalization, namely the *Fenchel–Young inequality*: for any conjugate pair  $(\psi, \psi^*)$  of proper, convex, lower-semicontinuous functions, one has  $\psi(s) + \psi^*(\xi) \geq \langle \xi, s \rangle$ , with the three equivalent characterizations of the equality case

$$\xi \in \partial\psi(s) \iff \psi(s) + \psi^*(\xi) = \langle \xi, s \rangle \iff s \in \partial\psi^*(\xi).$$

### The generalized step

Replace the two quadratic terms in (1.2) by a pair of convex *dissipation potentials*  $\mathcal{R}(x, \dot{x})$  and its Legendre dual  $\mathcal{R}^*(x, \xi)$ . The evolution equation becomes

$$\dot{x}(t) = D_\xi \mathcal{R}^*(x(t), -\nabla F(x(t))),$$

and, upon integration, the EDB becomes

$$\mathcal{L}(x, \dot{x}; [s, t]) := \int_s^t [\mathcal{R}(x, \dot{x}) + \mathcal{R}^*(x, -\nabla F(x))] \, dr + F(x(t)) - F(x(s)) = 0,$$

which makes sense whenever  $\mathcal{R}(x, \cdot)$  is convex with  $\mathcal{R}(x, 0) = 0$ . A *generalized gradient system* is the triple  $(\mathcal{X}, F, \mathcal{R})$  (equivalently  $(\mathcal{X}, F, \mathcal{R}^*)$ ).

When  $\mathcal{R}(x, v) = \frac{1}{2} |v|^2$  we recover the classical case.

## 1.3 Reversible Markov chains as gradient flows

sec:markov

Let  $V$  be a finite set of  $n \in \mathbb{N}$  *vertices*, and  $E := V \times V$  the set of *edges*. The state space is the simplex of probability vectors

$$\mathcal{P}(V) := \left\{ \rho: V \rightarrow [0, 1] \mid \rho_x \geq 0, \sum_x \rho_x = 1 \right\} \subset \mathbb{R}^n.$$

A continuous-time Markov chain  $(X_t)_{t \geq 0}$  has a *jump kernel*  $\kappa = (\kappa_{xy})$ , the rate of jumps from  $x$  to  $y$ , and generator

$$(Q\varphi)(x) = \sum_y [\varphi(y) - \varphi(x)] \kappa_{xy}, \quad \varphi \in \mathcal{B}(V).$$

Its law  $\rho(t)$  evolves by the *forward Kolmogorov equation*

$$\dot{\rho}(t) = Q^* \rho(t), \quad (Q^* \rho)_x = \sum_y [\rho_y \kappa_{yx} - \rho_x \kappa_{xy}], \quad \rho(0) = \bar{\rho} \in \mathcal{P}(V). \quad (1.3)$$

eq:fke-discrete

**Definition 1.4** (Detailed balance)  $\kappa$  satisfies *detailed balance* with respect to  $\pi \in \mathcal{P}(V)$  if the *edge measure*

def:db-discrete

$$\vartheta_{xy} := \pi_x \kappa_{xy} = \pi_y \kappa_{yx} = \vartheta_{yx} \quad \text{is symmetric.} \quad (DBC)$$

eq:db-discrete

Then  $\pi$  is stationary,  $Q^* \pi = 0$ .

Detailed balance is equivalent to time-reversibility of the chain started from  $\pi$ , and to self-adjointness of the semigroup  $T_t = e^{tQ}$  in  $L^2(\pi)$ ,  $\langle f, g \rangle_\pi := \sum_x f_x g_x \pi_x$ . Indeed,

$$\langle Qf, g \rangle_\pi = \sum_{x,y} \vartheta_{xy} f_y g_x = \sum_{x,y} \vartheta_{yx} f_x g_y = \langle f, Qg \rangle_\pi,$$

using symmetry of  $\vartheta$ ; reversibility follows by iterating.

### 1.3.1 The discrete continuity equation

Introduce the *graph gradient* and *graph divergence*

$$\begin{aligned} \nabla: \mathcal{B}(V) &\rightarrow \mathcal{B}(E), & (\nabla \varphi)(x, y) &= \varphi(y) - \varphi(x), \\ \operatorname{div}: \mathcal{M}(E) &\rightarrow \mathcal{M}(V), & (\operatorname{div} j)_x &= \sum_y (j_{xy} - j_{yx}), \end{aligned} \quad (1.4)$$

eq:graph-grad

which are negative adjoints:  $\sum_{x,y} \nabla \varphi(x, y) j_{xy} = - \sum_x \varphi(x) (\operatorname{div} j)_x$ . Defining the *net flux*  $j_{xy} := \frac{1}{2} (\rho_x \kappa_{xy} - \rho_y \kappa_{yx})$ —the half because each edge is counted in both orientations—a direct computation shows, for any test function  $\varphi$ ,

$$\sum_x \varphi(x) \dot{\rho}_x = \sum_{x,y} \rho_x \kappa_{xy} \nabla \varphi(x, y) = \sum_{x,y} \nabla \varphi(x, y) j_{xy} = - \sum_x \varphi(x) (\operatorname{div} j)_x,$$

the middle equality because  $\nabla\varphi$  is antisymmetric. Hence (1.3) is equivalent to the *continuity equation*

$$\begin{aligned}\dot{\rho} + \operatorname{div} j &= 0, \\ j_{xy} &= \frac{1}{2}(\rho_x \kappa_{xy} - \rho_y \kappa_{yx}).\end{aligned}\tag{CE}$$

Writing  $\rho = u\pi$  (i.e.  $\rho_x = u(x)\pi_x$ ) and using detailed balance,

$$j_{xy} = \frac{1}{2}(u(x)\pi_x \kappa_{xy} - u(y)\pi_y \kappa_{yx}) = \frac{1}{2}(u(x) - u(y)) \vartheta_{xy} = -\frac{1}{2} \nabla u(x, y) \vartheta_{xy}, \quad \text{eq:kr-discrete} \tag{1.5}$$

i.e., the flux is proportional to the gradient of the density along edges and is the precursor to the *kinetic relation*. Next, we introduce the energy and dissipation potentials.

### 1.3.2 The driving energy and the dissipation potential

By Sanov's theorem, the natural driving energy for the empirical measure of  $N$  independent copies of the chain is the *relative entropy*

$$\mathcal{F}(\rho) = \operatorname{Ent}_{\Phi}(\rho|\pi) := \begin{cases} \sum_x \Phi\left(\frac{\rho_x}{\pi_x}\right) \pi_x = \sum_x \rho_x \log \frac{\rho_x}{\pi_x}, & \rho \ll \pi, \\ +\infty, & \text{otherwise,} \end{cases}$$

with  $\Phi(s) = s \log s - s + 1$ , whose Fréchet derivative is  $D\mathcal{F}(\rho) = \log u$ ,  $u = d\rho/d\pi$ . More importantly, in the discrete setting there is *no underlying spatial gradient*: we must use the graph gradient  $\nabla$  in (1.4). This leaves a lot of freedom in the choice of dissipation potential. Writing  $\rho = u\pi$ , we consider only dual dissipation potentials of the form

$$\mathcal{R}^*(\rho, \xi) = \frac{1}{2} \sum_{x,y} \psi^*(\xi(x, y)) \mathfrak{m}(u(x), u(y)) \vartheta_{xy},$$

where  $\psi^*: \mathbb{R} \rightarrow [0, \infty)$  is even, convex, and superlinear with  $\psi^*(0) = 0$ , and the *flux multiplier*  $\mathfrak{m}$  is a symmetric, concave, 1-homogeneous *mean*, i.e.  $\mathfrak{m}(u, u) = u$ .

**Exercise 1.1** Show that the corresponding primal potential  $\mathcal{R}$ , with  $(\psi, \psi^*)$  a conjugate pair and  $2j_{xy} = w(x, y)\vartheta_{xy}$ , is given by

$$\mathcal{R}(\rho, j) = \frac{1}{2} \sum_{x,y} \Upsilon(u(x), u(y), w(x, y)) \vartheta_{xy},$$

where

$$\Upsilon(u, v, w) := \begin{cases} \mathfrak{m}(u, v) \psi\left(\frac{w}{\mathfrak{m}(u, v)}\right) & \text{if } \mathfrak{m}(u, v) > 0, \\ 0 & \text{if } w = 0, \\ +\infty & \text{if } w \neq 0 \text{ and } \mathfrak{m}(u, v) = 0, \end{cases} \quad \text{eq:upsilon} \tag{1.6}$$

is 1-homogeneous, jointly convex and lower semicontinuous.

### The $(\mathcal{F}, \mathcal{R}, \mathcal{R}^*)$ triple for a Markov chain

The chain (1.3) is the generalized gradient flow of

$$\dot{\rho} + \operatorname{div} j = 0, \quad j = D_{\xi} \mathcal{R}^*(\rho, -\nabla D\mathcal{F}(\rho)),$$

that is, the flux is produced by the dual dissipation potential acting on the thermodynamic force  $-\nabla D\mathcal{F}(\rho) = -\nabla \log u$ . Solutions are exactly the null-minimizers of the EDB functional

$$\mathcal{L}(\rho, j; [0, T]) = \int_0^T \left[ \mathcal{R}(\rho_t, j_t) + \mathcal{R}^*(\rho_t, -\nabla \log u_t) \right] dt + \mathcal{F}(\rho_T) - \mathcal{F}(\rho_0) \geq 0.$$

### 1.3.3 Two canonical choices

**Example 1.5** (Quadratic (Maas–Mielke–Chow–Huang–Zhou)) ex:quadratic Take  $\psi^*(\xi) = \frac{1}{2}\xi^2$  and  $\mathbf{m} = \mathbf{m}_{\log}$  the *logarithmic mean*  $\mathbf{m}_{\log}(u, v) = \frac{u-v}{\log u - \log v}$ . Then

$$j_{xy} = D_{\xi} \mathcal{R}^*(\rho, -\nabla \log u)_{xy} = \frac{1}{2} (-\nabla \log u(x, y)) \mathbf{m}_{\log}(u(x), u(y)) \vartheta_{xy} = \frac{1}{2} (u(x) - u(y)) \vartheta_{xy},$$

the factor  $\frac{1}{2}$  coming from the one in  $\mathcal{R}^*$ , and  $(-\nabla \log u(x, y)) \mathbf{m}_{\log}(u(x), u(y)) = u(x) - u(y)$  cancelling the logarithmic mean against its own denominator. This is exactly (1.5). This is the structure of Maas [8], Mielke [9], and Chow–Huang–Li–Zhou [2]; it makes  $\mathcal{P}(V)$  a Riemannian-like metric space (a discrete Wasserstein space) and underlies the Erbar–Maas notion of discrete Ricci curvature [4].

**Example 1.6** (Cosh/large-deviation structure) ex:cosh Take

$$\psi^*(\xi) = 4(\cosh(\xi/2) - 1), \quad \mathbf{m}(a, b) = \sqrt{ab},$$

so that, with  $\rho = u\pi$ , the edge weight is the geometric mean  $\sqrt{u(x)u(y)} \vartheta_{xy}$ . Since  $(\psi^*)'(\xi) = 2 \sinh(\xi/2)$ , the flux relation reads

$$j_{xy} = \frac{1}{2} (\psi^*)'(-\nabla \log u(x, y)) \sqrt{u(x)u(y)} \vartheta_{xy} = \frac{1}{2} \left[ 2 \sinh\left(\frac{1}{2} \log \frac{u(x)}{u(y)}\right) \right] \sqrt{u(x)u(y)} \vartheta_{xy}.$$

Using  $2 \sinh\left(\frac{1}{2} \log \frac{u(x)}{u(y)}\right) = \sqrt{u(x)/u(y)} - \sqrt{u(y)/u(x)}$ , the bracket against  $\sqrt{u(x)u(y)}$  is  $u(x) - u(y)$ , so  $j_{xy} = \frac{1}{2} (u(x) - u(y)) \vartheta_{xy}$ , again recovering (1.5). The primal potential is

$$\psi(s) = 2s \log \left( \frac{s + \sqrt{s^2 + 4}}{2} \right) - 2\sqrt{s^2 + 4} + 4.$$

This structure is *not*  $p$ -homogeneous, hence *not* metric, but it is distinguished because it arises as the large-deviation rate functional of the empirical fluxes of the underlying chain [12]. It is this non-metric, large-deviation-induced structure that the rigorous theory of Lecture 2 is built to handle.

## 1.4 The Fisher information and dissipation along solutions

Setting  $s = 0$  and differentiating the EDB, a solution dissipates energy at rate

$$\frac{d}{dt}\mathcal{F}(\rho_t) = -\mathcal{R}(\rho_t, j_t) - \mathcal{D}(\rho_t) \leq 0, \quad \mathcal{D}(\rho) := \mathcal{R}^*(\rho, -\nabla \log u),$$

where  $\mathcal{D}$  is the (generalized) *Fisher information*. For the two canonical choices,

$$\begin{aligned} \mathcal{D}_{\text{quad}}(\rho) &= \frac{1}{4} \sum_{x,y} (\log u(x) - \log u(y))(u(x) - u(y)) \vartheta_{xy}, \\ \mathcal{D}_{\text{cosh}}(\rho) &= \sum_{x,y} (\sqrt{u(x)} - \sqrt{u(y)})^2 \vartheta_{xy}. \end{aligned}$$

Both are nonnegative, vanish only at  $u \equiv \text{const}$  (i.e.,  $\rho = \pi$ ), and serve as Lyapunov dissipation rates driving convergence to equilibrium.

## 2 Generalized gradient structures for jump processes

ch: jump

### Goal of Lecture 2

We lift the discrete picture of Lecture 1 to jump processes on a general *Polish* state space  $V$ , where the set of vertices may be finite, countable, or uncountable. We make the cosh-type structure rigorous: the continuity equation on  $\mathcal{M}^+(V)$ , the dissipation potential  $\mathcal{R}$ , the chain rule, the Fisher information  $\mathcal{D}$ , and the  $(\mathcal{F}, \mathcal{R}, \mathcal{R}^*)$  energy–dissipation balance that *defines* solutions. We sketch existence via minimizing movements. This is the framework of Peletier–Rossi–Savaré–Tse [14].

Think of a jump process as moving from a *vertex*  $x$  to a *vertex*  $y$  along an *edge*  $(x, y)$  of a *graph*. Now  $V$  is a *Polish space*—separable and completely metrizable—carrying its Borel  $\sigma$ -algebra  $\mathcal{B}$ , and  $E := V \times V$ . The law of the process is a time-dependent measure  $t \mapsto \rho_t \in \mathcal{M}^+(V)$  solving the *forward Kolmogorov equation*

$$\partial_t \rho_t = Q^* \rho_t, \quad (Q^* \rho)(dx) = \int_V \rho(dy) \kappa(y, dx) - \rho(dx) \int_V \kappa(x, dy), \quad \text{eq: fke-jump} \quad (2.1)$$

where  $Q^*$  is the dual of the generator

$$(Q\varphi)(x) = \int_V [\varphi(y) - \varphi(x)] \kappa(x, dy), \quad \varphi \in \mathcal{B}_b(V).$$

The *jump kernel*  $\kappa(x, \cdot) \in \mathcal{M}^+(V)$  is the infinitesimal rate of jumps from  $x \in V$ .

ass: Vpik

**Assumption 2.1**  $V$  is Polish with Borel  $\sigma$ -algebra  $\mathcal{B}$ , the reference measure is  $\pi \in \mathcal{M}^+(V)$ , and  $(\kappa(x, \cdot))_{x \in V}$  is a kernel satisfying  $\sup_x \kappa(x, V) < +\infty$  and the *detailed-balance condition* w.r.t.  $\pi$ , i.e.,

$$\int_A \kappa(x, B) \pi(dx) = \int_B \kappa(y, A) \pi(dy) \quad \text{for all } A, B \in \mathcal{B}. \quad \text{eq: db-jump} \quad (\text{DBC})$$

Detailed balance (DBC) says exactly that the *edge measure*

$$\vartheta(dx dy) := \pi(dx) \kappa(x, dy) \in \mathcal{M}^+(E) \quad \text{is symmetric.}$$

This is the continuum analog of the discrete condition in Lecture 1.

**Graph gradient and divergence.** These carry over directly from §1.3:

$$\begin{aligned} (\nabla \varphi)(x, y) &:= \varphi(y) - \varphi(x), \quad \varphi \in \mathcal{B}_b(V) \\ (\operatorname{div} j)(dx) &:= \int_{y \in V} [j(dx, dy) - j(dy, dx)], \quad j \in \mathcal{M}(E). \end{aligned}$$

As before, they satisfy the integration-by-parts formula  $\langle \nabla \varphi, j \rangle = -\langle \varphi, \operatorname{div} j \rangle$ .

**Setwise convergence.** The ‘correct’ notion of convergence for sequences of measures in this setting is that of *setwise convergence*. It sits strictly between convergence in total variation and narrow (or weak) convergence. Every limit below is taken in this sense.

**Definition 2.2** (Setwise convergence) A sequence  $(\mu_n)_n \subset \mathcal{M}(V)$  converges *setwise* to  $\mu \in \mathcal{M}(V)$  if def:setwise

$$\mu_n(B) \longrightarrow \mu(B) \quad \text{for every } B \in \mathcal{B}.$$

The same definition on  $\mathcal{M}(E)$  is used for fluxes.

**Proposition 2.3.** *For sequences  $(\mu_n)_n \subset \mathcal{M}(V)$  and  $\mu \in \mathcal{M}(V)$ , the following characterizations of setwise convergence are equivalent:*

- (i)  $\mu_n \rightarrow \mu$  setwise in  $\mathcal{M}(V)$ .
- (ii) (Convergence in duality)  $\langle \varphi, \mu_n \rangle \rightarrow \langle \varphi, \mu \rangle$  for every  $\varphi \in \mathcal{B}_b(V)$ .
- (iii) (Weak  $L^1$ -convergence of densities) There exists a common dominating measure  $\gamma \in \mathcal{M}(V)$  for  $(\mu_n)_n$  and  $\mu$ , i.e.,  $\mu_n \ll \gamma$ ,  $\mu \ll \gamma$ , such that

$$\frac{d\mu_n}{d\gamma} \rightharpoonup \frac{d\mu}{d\gamma} \quad \text{weakly in } L^1(V, \gamma).$$

## 2.1 The three ingredients of the gradient structure

### 2.1.1 Driving energy

**Assumption 2.4** ( $\mathcal{F}\Phi$ ) The driving functional is the  $\Phi$ -relative entropy ass:Ephi

$$\mathcal{F}(\rho) = \text{Ent}_\Phi(\rho|\pi) := \begin{cases} \int_V \Phi\left(\frac{d\rho}{d\pi}\right) d\pi, & \rho \ll \pi, \\ +\infty, & \text{otherwise,} \end{cases}$$

where  $\Phi \in C([0, \infty)) \cap C^1((0, \infty))$  is convex with superlinear growth and  $\min \Phi = 0$ . The central example is the Boltzmann–Shannon density  $\Phi_B(s) = s \log s - s + 1$ .

$\mathcal{F}$  is lower semicontinuous for setwise convergence, and its sublevels are setwise relatively compact—the compactness backbone of the existence theory.

### 2.1.2 The continuity equation on $\mathcal{M}^+(V)$

**Definition 2.5** (Continuity equation) Let  $I = [a, b]$ . A pair  $(\rho, j)$  belongs to  $\text{CE}(I)$  if def:ce-jump

- (i)  $(\rho_t)_{t \in I} \in \mathcal{M}^+(V)$  is a family of time-dependent measures,
- (ii)  $(j_t)_{t \in I} \subset \mathcal{M}(E)$  is a measurable family with  $\int_I |j_t|(E) dt < +\infty$ , and
- (iii) the continuity equation  $\partial_t \rho_t + \text{div } j_t = 0$  is satisfied in the following sense:

$$\int_V \varphi d\rho_{t_2} - \int_V \varphi d\rho_{t_1} = \int_{t_1}^{t_2} \int_E \nabla \varphi \cdot dj_t dt \quad \forall \varphi \in \mathcal{B}_b(V), [t_1, t_2] \subseteq I, \quad (\text{CE})$$

We write  $\text{CE}(I; \rho_a)/\text{CE}(I; \rho_a, \rho_b)$  when the endpoints are prescribed.

lem:CE-implications

**Lemma 2.6.** *If  $(\rho, j) \in \text{CE}(I)$ , then there is a common dominating measure  $\gamma \in \mathcal{M}^+(V)$  for  $(\rho_t)_{t \in I}$  (i.e.,  $\rho_t \ll \gamma$  for all  $t \in I$ ), and an absolutely continuous map  $u : I \rightarrow L^1(V, \gamma)$  such that  $\rho_t = u_t \gamma \ll \gamma$  for every  $t \in I$ .*

*In particular,  $t \mapsto \rho_t$  is continuous w.r.t. the total variation distance, and taking  $\varphi \equiv 1$  shows that the total mass  $\rho_t(V)$  is conserved in time.*

### 2.1.3 The dissipation potential

ass:Rstalarpha

**Assumption 2.7** The dual dissipation density  $\Psi^* : \mathbb{R} \rightarrow [0, \infty)$  is even, convex, differentiable, with  $\Psi^*(0) = 0$  and superlinear growth. The *flux density*  $\mathbf{m} : [0, +\infty)^2 \rightarrow [0, +\infty)$  is continuous, symmetric, concave and 1-homogeneous, and is a *mean*,  $\mathbf{m}(u, u) = u$ .

Concavity and 1-homogeneity make  $\mathbf{m}$  superadditive, so for  $u \leq v$  we get  $\mathbf{m}(u, v) \geq \mathbf{m}(u, u) + \mathbf{m}(0, v - u) \geq u$ : the mean dominates  $\min(u, v)$ , and in particular  $\mathbf{m}(u, v) > 0$  whenever  $uv > 0$ . Along a curve  $\rho = u\pi$ , the *dual dissipation potential* is

$$\mathcal{R}^*(\rho, \xi) = \frac{1}{2} \int_E \Psi^*(\xi(x, y)) \mathbf{m}(u(x), u(y)) \vartheta(dx dy), \quad \xi \in \mathcal{B}_b(E).$$

Its Legendre dual, the *primal* potential, has the perspective-function form: writing  $2j = w \vartheta$  (so  $w$  is a density of the flux against the symmetric edge measure),

$$\mathcal{R}(\rho, j) = \frac{1}{2} \int_E \Upsilon(u(x), u(y), w(x, y)) \vartheta(dx dy),$$

where  $\Upsilon$ , defined as in (1.6), is jointly convex and lower semicontinuous.

#### The cosh structure on $\mathcal{M}^+(V)$

The distinguished choice, induced by the large deviations of the underlying particle system (Lecture 3), is

$$\Psi^*(\xi) = 4(\cosh(\xi/2) - 1), \quad \mathbf{m}(u, v) = \sqrt{uv},$$

for which the formal gradient-flow equation  $\partial_t \rho = -\text{div } D_\xi \mathcal{R}^*(\rho, -\nabla \Phi'(u))$  reduces, with  $\Phi'(u) = \log u$ , to the linear forward Kolmogorov equation (2.1)

$$\partial_t u_t(x) = \int_V (u_t(y) - u_t(x)) \kappa(x, dy).$$

**Remark 2.8** (Linear compatibility and the general equation) For a triple  $(\Psi^*, \Phi, \mathbf{m})$ , equation  $\partial_t \rho = -\text{div } D_\xi \mathcal{R}^*(\rho, -\nabla \Phi'(u))$  becomes the integro-differential equation

$$\partial_t u_t(x) = \int_V F(u_t(x), u_t(y)) \kappa(x, dy)$$

with  $F(u, v) = (\Psi^*)'(\Phi'(v) - \Phi'(u))\mathbf{m}(u, v)$ . It coincides with the *linear* equation (2.1) precisely when the *linear compatibility condition*  $F(u, v) = v - u$  holds for all  $u, v > 0$ . Both the *cosh* choice and the quadratic choice  $\Psi^*(\xi) = \frac{1}{2}\xi^2$ ,  $\mathbf{m}(u, v) = \frac{u-v}{\log u - \log v}$  satisfy it.

## 2.2 The Fisher information

We now define the *Fisher information*. For this, we will need the following objects.

**Definition 2.9** (Force and Fisher integrands) def:fisher-integrands Following [14], set  $\phi'(0) := \lim_{r \downarrow 0} \phi'(r) \in [-\infty, +\infty)$  and define the *discrete force*

$$A_\phi(u, v) := \begin{cases} \phi'(v) - \phi'(u), & (u, v) \in \mathbb{R}_+^2 \setminus \{(0, 0)\}, \\ 0, & u = v = 0. \end{cases}$$

For the Boltzmann density, this is  $A_\phi(u, v) = \log(v/u)$ , with values  $\pm\infty$  on the axes. Set

$$D_\phi^-(u, v) := \begin{cases} \psi^*(A_\phi(u, v))\mathfrak{m}(u, v), & \mathfrak{m}(u, v) > 0, \\ 0, & \text{otherwise,} \end{cases}$$

$$D_\phi^+(u, v) := \begin{cases} \psi^*(A_\phi(u, v))\mathfrak{m}(u, v), & \mathfrak{m}(u, v) > 0, \\ 0, & u = v = 0, \\ +\infty, & \mathfrak{m}(u, v) = 0, \ u \neq v. \end{cases}$$

The *Fisher information* is then defined as

$$\text{dom}(\mathcal{F}) \ni \rho \mapsto \mathcal{D}(\rho) = \frac{1}{2} \int_E D_\phi(u(x), u(y)) \vartheta(dx dy), \quad \rho = u\pi,$$

where  $D_\phi$  is the lower-semicontinuous envelope of  $D_\phi^+$  on  $\mathbb{R}_+^2$ , i.e.,

$$D_\phi(u, v) := \sup \left\{ G(u, v) : G : \mathbb{R}_+^2 \rightarrow [0, +\infty] \text{ lower semicontinuous, } G \leq D_\phi^+ \right\}$$

$$= \inf \left\{ \liminf_{n \rightarrow \infty} D_\phi^+(u_n, v_n) : u_n \rightarrow u, \ v_n \rightarrow v \right\}.$$

The integrands  $D_\phi^\pm$  and  $D_\phi$  differ only on the degenerate set  $\{\mathfrak{m} = 0, \ u \neq v\}$  and always satisfy  $D_\phi^- \leq D_\phi \leq D_\phi^+$ .

**Proposition 2.10** (Lower semicontinuity of  $\mathcal{D}$ ). prop:fisher-lsc Assume either  $\pi$  is purely atomic or that  $D_\phi$  is convex on  $\mathbb{R}_+^2$ . Then  $\mathcal{D}$  is lower semicontinuous w.r.t. setwise convergence (Definition 2.2), i.e., for all  $(\mu_n)_n \subset \text{dom}(\mathcal{F})$ ,  $\mu \in \text{dom}(\mathcal{F})$ , one has that

$$\mu_n \rightarrow \mu \text{ setwise in } \mathcal{M}^+(V) \implies \mathcal{D}(\mu) \leq \liminf_{n \rightarrow \infty} \mathcal{D}(\mu_n).$$

**Example 2.11** ex:fisher For the linear equation with  $\phi = \phi_B$ , so  $A_\phi(u, v) = \log(v/u)$ , we find

(1) if  $\psi^*(s) = \frac{1}{2}s^2$  and  $\mathfrak{m}(u, v) = \frac{u-v}{\log u - \log v}$ , then

$$D_\phi(u, v) = \begin{cases} \frac{1}{2}(\log u - \log v)(u - v) & \text{for } uv > 0, \\ 0 & \text{for } u = v = 0, \\ +\infty & \text{for } uv = 0, \ u \neq v; \end{cases}$$

(2) if  $\psi^*(s) = 4(\cosh \frac{s}{2} - 1)$  and  $\mathfrak{m}(u, v) = \sqrt{uv}$ , then

$$D_\Phi(u, v) = 2(\sqrt{u} - \sqrt{v})^2 \quad \forall (u, v) \in [0, +\infty) \times [0, +\infty).$$

Both examples give convex  $D_\Phi$  on  $\mathbb{R}_+^2$ , so Proposition 2.10 applies and  $\mathcal{D}$  is setwise lower semicontinuous in either case. The difference between them occurs on  $\{uv = 0, u \neq v\}$ . We note that finiteness of the Fisher information  $\mathcal{D}$  is trivial for the *cosh* example, but not for the quadratic case. On the other hand, if we know that the  $\mathcal{D}$  is finite for the quadratic case, then we know that  $u(x)u(y) > 0$  for  $\vartheta$ -a.e.  $(x, y) \in E$ , and hence,  $u(x) > 0$  for every  $\pi$ -a.e.  $x \in V$ , while no such property can be claimed for (2).

## 2.3 Admissible curves and the chain rule

The heart of the rigorous theory is a *chain rule*, which we show to hold for the following class of curves.

def:admissible

**Definition 2.12** (Admissible curves) A pair  $(\rho, j) \in \text{CE}(I)$  is *admissible* if

$$\sup_{t \in I} \mathcal{F}(\rho_t) < +\infty, \quad \int_I [\mathcal{R}(\rho_t, j_t) + \mathcal{D}(\rho_t)] dt < +\infty.$$

In this case, we write  $(\rho, j) \in \mathcal{A}(I)$ .

The three terms in the definition are exactly the terms of the balance (EDB), so  $\mathcal{A}(I)$  is simply the class on which (EDB) is a statement about finite quantities.

A few remarks are in order:

**Remark 2.13** (1) The finiteness of  $\mathcal{F}(\rho_t)$  gives  $\rho_t = u_t \pi$  with  $u_t \in L^1(V, \pi)$ . Hence,  $(\rho, j) \in \mathcal{A}(I) \subset \text{CE}(I)$  implies  $\rho_t = u_t \pi$  with  $u_t \in L^1(V, \pi)$  for all  $t \in I$ .

(2) The finiteness of  $\mathcal{R}(\rho_t, j_t)$  forces  $j_t$  to vanish on  $\{\mathfrak{m} = 0\}$  for almost every  $t \in I$ , since  $\Upsilon(u, v, w) = +\infty$  for  $w \neq 0$  there. On  $\{\mathfrak{m} > 0\}$ , we have that  $D_\Phi^\pm = D_\Phi$ .

Under finite action, we obtain

**Proposition 2.14.** Let  $(\rho, j) \in \text{CE}(I)$  with finite action  $\int_I \mathcal{R}(\rho_t, j_t) dt < +\infty$ . If  $\rho_t \ll \pi$  for some  $t \in I$ , then there exists

- (1) an absolutely continuous and a.e. differentiable map  $u : I \rightarrow L^1(V, \pi)$  with  $\rho_t = u_t \pi$ ,
- (2) a map  $w \in L^1(I \times E, \lambda \otimes \vartheta)$  such that  $2j_t = w_t \vartheta$  for a.e.  $t \in I$ , and
- (3) the equation

$$\partial_t u_t = -\frac{1}{2} \int_V [w_t(x, y) - w_t(y, x)] \kappa(x, dy)$$

is satisfied for almost every  $t \in I$ .

We conclude this subsection with the chain rule for admissible pairs  $(\rho, j) \in \mathcal{A}(I)$ .

thm:chainrule

**Theorem 2.15** (Chain rule [14]). Let  $(\rho, j) \in \mathcal{A}(I)$  and write  $\rho_t = u_t \pi$ . Then  $t \mapsto \mathcal{F}(\rho_t)$  is absolutely continuous on  $I$  and, for all  $s \leq t$  in  $I$ ,

$$\mathcal{F}(\rho_t) - \mathcal{F}(\rho_s) = \int_s^t \int_E \nabla \Phi'(u_r) dj_r dr \geq - \int_s^t \mathcal{R}(\rho_r, j_r) + \mathcal{D}(\rho_r) dr. \quad (\text{CR})$$

eq:chainrule

## 2.4 Definition of solutions: the $(\mathcal{F}, \mathcal{R}, \mathcal{R}^*)$ balance

**Definition 2.16** ( $(\mathcal{F}, \mathcal{R}, \mathcal{R}^*)$  energy–dissipation balance) A pair  $(\rho, j) \in \text{CE}(0, T)$  is a *solution* of the  $(\mathcal{F}, \mathcal{R}, \mathcal{R}^*)$  system if for every  $0 \leq s \leq t \leq T$ ,

$$\mathcal{L}(\rho, j; [s, t]) := \int_s^t [\mathcal{R}(\rho_r, j_r) + \mathcal{D}(\rho_r)] dr + \mathcal{F}(\rho_t) - \mathcal{F}(\rho_s) = 0. \quad (\text{EDB})$$

### EDB = EDI under the chain rule

The chain rule (CR) and Young’s inequality  $\mathcal{R}(\rho, j) + \mathcal{R}^*(\rho, -\nabla \phi'(u)) + \langle \nabla \phi'(u), j \rangle \geq 0$  yields, for every  $(\rho, j) \in \text{CE}(0, T)$ , the *chain-rule (lower) bound*

$$\mathcal{L}(\rho, j; [s, t]) \geq 0.$$

Hence, the balance (EDB) is equivalent to the energy–dissipation *inequality* (EDI)

$$\mathcal{L}(\rho, j; [s, t]) \leq 0. \quad (\text{EDI})$$

**Theorem 2.17** (Equation form of solutions [14]). A curve  $(\rho_t)$  with  $\mathcal{F}(\rho_0) < +\infty$  solves the  $(\mathcal{F}, \mathcal{R}, \mathcal{R}^*)$  system iff

- (i)  $\rho_t = u_t \pi$  with  $t \mapsto u_t$  absolutely continuous in  $L^1(V, \pi)$ ,
- (ii) the equality case  $D_\phi = D_\phi^-$  holds  $\lambda \otimes \vartheta$ -a.e., and
- (iii) the flux takes the form  $2j_t(dx dy) = -F_0(u_t(x), u_t(y)) \vartheta(dx dy)$  such that

$$\partial_t u_t(x) = \int_V F_0(u_t(x), u_t(y)) \kappa(x, dy) \quad (\lambda \otimes \pi)\text{-a.e.},$$

where  $F_0(u, v) = (\psi^*)'(A_\phi(u, v))\mathfrak{m}(u, v)$  is the skew-symmetric flux map.

**Theorem 2.18** (Uniqueness [14]). If the Fisher information  $\mathcal{D}$  is convex and  $\phi$  is strictly convex, then solutions of (EDB) are uniquely determined by  $\rho_0 \in \text{dom}(\mathcal{F})$ .

## 2.5 Existence via minimizing movements

Solutions in the sense of Definition 2.16 are produced by the *Minimizing-Movement* (or JKO) scheme. Fix a time step  $\tau > 0$  and let  $\rho_\tau^0 = \rho_0 \in \text{dom}(\mathcal{F})$ . We then recursively set

$$\rho_\tau^n \in \arg \min_{\rho \in \mathcal{M}^+(V)} \left\{ \mathcal{W}_\tau(\rho_\tau^{n-1}, \rho) + \mathcal{F}(\rho) \right\}, \quad (2.2)$$

where the *Dynamical-Variational Transport* cost  $\mathcal{W}_\tau: \mathcal{M}^+(V)^2 \rightarrow [0, +\infty)$  is

$$\mathcal{W}_\tau(\mu, \nu) := \inf \left\{ \int_0^\tau \mathcal{R}(\rho_r, j_r) dr : (\rho, j) \in \text{CE}(0, \tau; \mu, \nu) \right\}.$$

**Theorem 2.19** (Convergence of the scheme [14]). Let  $\mathcal{F}(\rho_0) < +\infty$ . Under Assumptions 2.1, 2.4, 2.7, the piecewise-constant interpolants of (2.2) admit, as  $\tau \rightarrow 0$ , a (sub)sequential limit  $(\rho, j) \in \text{CE}(0, T)$  satisfying the energy–dissipation inequality (EDI). In particular  $(\rho, j) \in \mathcal{A}(0, T)$ , which with the chain rule (CR) makes it a solution of the  $(\mathcal{F}, \mathcal{R}, \mathcal{R}^*)$  system in the sense of Definition 2.16.

## 3 Where gradient structures come from

ch:origins

### Goal of Lecture 3

Lecture 1 exhibited a gradient structure for a reversible Markov chain and remarked, without justification, that the cosh structure is “singled out by large deviations”. This lecture makes good on that promise. Its message is the *selection principle*: for a stochastic particle system the large-deviation rate functional is canonical; three classical theorems—Sanov, Girsanov, and disintegration—turn it into an explicit action on curves of measures and fluxes; and splitting that action into its symmetric and antisymmetric parts under *time reversal* produces a gradient structure, with the symmetric part the dissipation and the antisymmetric part the change in free energy. Detailed balance is exactly the condition making this work. Diffusions come out quadratic, jump processes come out cosh.

### 3.1 The selection principle: large deviations

sec:selection

Lecture 1 left a question open. The reversible Markov chain admits at least two gradient structures—the quadratic one of Example 1.5 and the cosh one of Example 1.6—and both generate the same forward Kolmogorov equation. Which should one choose? Answering “whichever is convenient” is unsatisfying, because different structures predict different behavior away from equilibrium and behave differently under limits. We want a principled answer, and probability gives us one.

The key point is that the macroscopic equation isn’t fundamental. We can derive it from the law of large numbers for a stochastic particle system. This system captures not only the average behavior but also the *probability of deviations* from that average. This information is canonical—free of modeling choices—and also reveals a gradient structure.

The derivation has three steps, each of which is a classical theorem. We state it for  $N \in \mathbb{N}$  independent copies of a Markov jump process with jump kernel  $\kappa$  on a Polish space  $V$ , writing  $E := V \times V$  as in Lecture 1, started from measure  $\bar{\rho}$ . Let  $\mathbf{P}_{\bar{\rho}}$  denote the law of a single ‘particle’, which is a probability measure on the space of *càdlàg* paths  $D([0, T]; V)$ , i.e., paths  $\omega : [0, T] \rightarrow V$  that are right-continuous with left limits, equipped with the Skorohod  $J_1$  topology. Since  $V$  is Polish, so is  $D([0, T]; V)$ , w.r.t. the  $J_1$  topology.

**Step 1: Sanov, on path space.** Consider the *empirical process* on path space

$$\mathbf{P}^N := \frac{1}{N} \sum_{i=1}^N \delta_{X^i} \in \mathcal{P}(D([0, T]; V)).$$

This is a random measure not on states but on whole trajectories. Since the trajectories are independent, Sanov’s theorem applies, and we obtain a large-deviation principle with speed  $N$  whose rate function is a relative entropy on path space given by

$$\text{Prob}(\mathbf{P}^N \approx \mathbf{Q}) \asymp e^{-N J_{LD}(\mathbf{Q})}, \quad J_{LD}(\mathbf{Q}) = \text{Ent}(\mathbf{Q} | \mathbf{P}_{\bar{\rho}}). \quad \text{eq:sanov-path} \quad (3.1)$$

This is the full probabilistic input, and it offers no modeling freedom. Note that it vanishes exactly at  $\mathbf{Q} = \mathbf{P}_{\bar{\rho}}$ , the law of the typical trajectory.

**Step 2: Girsanov for jump processes.** Formula (3.1) is not useful in its current form because it represents an entropy between measures on an infinite-dimensional path space. The relevant tool, analogous to Girsanov's theorem for diffusions but for jump processes, precisely determines which path measures  $\mathbf{Q}$  have finite entropy relative to  $\mathbf{P}_{\bar{\rho}}$ , and calculates their density.

Consider the law  $\mathbf{Q}^{\hat{\kappa}}$  of the process which starts from  $\rho_0$  and jumps with a *tilted*, possibly time-dependent kernel

$$\hat{\kappa}_t(x, dy) = r_t(x, y) \kappa(x, dy), \quad \text{eq:tilted-kernel} \quad (3.2)$$

so that the density  $r_t$  records, at each state and each target, the factor by which *observed* rates depart from the *expected* ones. Then  $\mathbf{Q}^{\hat{\kappa}} \ll \mathbf{P}_{\bar{\rho}}$ , with density given by the stochastic (Doléans-Dade) exponential: for a càdlàg path  $\omega = (\omega_t)_{t \in [0, T]} \in D([0, T]; V)$ ,

$$\frac{d\mathbf{Q}^{\hat{\kappa}}}{d\mathbf{P}_{\bar{\rho}}}(\omega) = \frac{d\rho_0}{d\bar{\rho}}(\omega_0) \exp\left(\int_0^T \int_V (1 - r_t(\omega_{t-}, y)) \kappa(\omega_{t-}, dy) dt\right) \prod_{\substack{t \leq T: \\ \omega_{t-} \neq \omega_t}} r_t(\omega_{t-}, \omega_t), \quad \text{eq:girsanov-dd} \quad (3.3)$$

one factor  $r_t$  per *observed* jump, times an exponential correction accounting for the jumps that were *not* observed. For each càdlàg path  $\omega$  we denote by

$$\mathbf{N}_{\omega} := \sum_{t \leq T: \omega_{t-} \neq \omega_t} \delta_{(t, \omega_t)} \in \mathcal{M}((0, T] \times V),$$

its associated *jump measure*, which records when the path jumps and where it lands. Taking the logarithm of (3.3) turns the product into an integral against  $\mathbf{N}_{\omega}$ , and the resulting identity splits into an initial term and a dynamic one:

$$\begin{aligned} \log \frac{d\mathbf{Q}^{\hat{\kappa}}}{d\mathbf{P}_{\bar{\rho}}}(\omega) &= \underbrace{\log \frac{d\rho_0}{d\bar{\rho}}(\omega_0)}_{\text{initial data}} \\ &+ \underbrace{\int_0^T \int_V \log r_t(\omega_{t-}, y) \mathbf{N}_{\omega}(dt dy) - \int_0^T \int_V (r_t(\omega_{t-}, y) - 1) \kappa(\omega_{t-}, dy) dt}_{\text{dynamic part}}. \end{aligned} \quad \text{eq:girsanov-jump} \quad (3.4)$$

The dynamic part has a transparent structure: the first integral pays  $\log r_t$  for each jump the path *actually makes*, while the second accounts for the jumps that were expected but *not observed*. The rigorous statement follows from the *martingale problem*. We note that the jump measure  $\mathbf{N}_{\omega}$  has compensator  $\hat{\kappa}_t(\omega_t, dy) dt$  under  $\mathbf{Q}^{\hat{\kappa}}$ , i.e., for every nonnegative predictable function  $f$ , we have that

$$\mathbb{E}_{\mathbf{Q}^{\hat{\kappa}}} \left[ \int_0^T \int_V f(t, \omega_{t-}, y) \mathbf{N}_{\omega}(dt dy) \right] = \mathbb{E}_{\mathbf{Q}^{\hat{\kappa}}} \left[ \int_0^T \int_V f(t, \omega_{t-}, y) \hat{\kappa}_t(\omega_{t-}, dy) dt \right]. \quad \text{eq:compensator} \quad (3.5)$$

Now write  $\rho_t := (e_t)_{\#} \mathbf{Q}^{\hat{\kappa}}$  for the time marginal along the map  $e_t(\omega) = \omega_t$ , and set

$$\text{Ent}(\hat{\kappa}_t(x, \cdot) | \kappa(x, \cdot)) := \int_V \Phi(r_t(x, y)) \kappa(x, dy), \quad \Phi(s) = s \log s - s + 1,$$

By applying (3.5) with  $f = \log r_t$ , the jump-measure integral in (3.4) turns into a time integral against  $\hat{\kappa}_t = r_t \kappa$ . Together with the fact that

$$\mathbb{E}_{\mathbf{Q}^{\hat{\kappa}}} \left[ \int_0^T g(t, \omega_{t-}) dt \right] = \int_0^T \int_V g(t, x) \rho_t(dx),$$

for every predictable function  $g$ , the dynamic part in (3.4) combine to a single  $\Phi$ -integral,

$$\int_0^T \int_E \Phi(r_t(x, y)) \kappa(x, dy) \rho_t(dx) dt = \int_0^T \int_V \text{Ent}(\hat{\kappa}_t(x, \cdot) | \kappa(x, \cdot)) \rho_t(dx) dt, \quad \text{eq:dynamic-average} \quad (3.6)$$

and the initial term of (3.4) gives  $\text{Ent}(\rho_0 | \bar{\rho})$ . Since both terms are nonnegative,

$$\text{Ent}(\mathbf{Q}^{\hat{\kappa}} | \mathbf{P}_{\bar{\rho}}) < \infty \iff \begin{cases} \text{Ent}(\rho_0 | \bar{\rho}) < +\infty & \text{and} \\ \int_0^T \int_V \text{Ent}(\hat{\kappa}_t(x, \cdot) | \kappa(x, \cdot)) \rho_t(dx) dt < +\infty. \end{cases}$$

Both conditions hold, for instance, when the initial entropy is finite, the tilt is bounded,  $r_t \leq C$ , and the total rates are bounded,  $\sup_x \kappa(x, V) < +\infty$ —these are the standing assumptions of Lecture 2. Note the two regimes of  $\Phi$ : near  $s = 1$  one has  $\Phi(s) \approx \frac{1}{2}(s-1)^2$ , so mildly tilted rates cost quadratically, while  $\Phi(s) \sim s \log s$  as  $s \rightarrow +\infty$ , so large tilts pay a Boltzmann price—a first shadow of exponential-type dissipation.

The direction that matters for us is the converse, established in great generality by Conforti and Léonard [3]: if  $\text{Ent}(\mathbf{Q} | \mathbf{P}_{\bar{\rho}}) < \infty$ , then  $\mathbf{Q}$  is itself the law of a jump process with the *same* deterministic drift and a *predictable* jump intensity kernel  $\hat{\kappa}_t \ll \kappa$ . In other words, every finite-entropy path measure is of the tilted form (3.3). Finiteness of the entropy is thus not merely an integrability condition, but also a structural statement. It says that an atypical trajectory of the particle system is simply the same process with tilted rates, and  $\hat{\kappa}_t$  is precisely the observed jump kernel.

rem:markovianization

**Remark 3.1** (From predictable to Markov) A priori, the intensity produced by the converse is predictable, so  $\mathbf{Q}$  need not be Markov. However, if the reference process  $\mathbf{P}_{\bar{\rho}}$  is Markov, then we can assume that the tilt is Markovian (see [3]).

**Step 3: From kernels to fluxes.** Collecting the averages of Step 2, we obtain

$$\text{Ent}(\mathbf{Q} | \mathbf{P}_{\bar{\rho}}) = \underbrace{\text{Ent}(\rho_0 | \bar{\rho})}_{\text{initial data}} + \underbrace{\int_0^T \int_V \text{Ent}(\hat{\kappa}_t(x, \cdot) | \kappa(x, \cdot)) \rho_t(dx) dt}_{\text{dynamic cost}}, \quad \text{eq:disintegrate-entropy} \quad (3.7)$$

Next, we eliminate the kernel in favor of the mass it transports. Define the *flux* and the *expected flux* at a state  $\rho$  by

$$j_t := \rho_t \otimes \hat{\kappa}_t, \quad (\rho \otimes \kappa)(dx dy) := \rho(dx) \kappa(x, dy),$$

both measures on  $E$ . Relative entropy disintegrates over the first variable, so that

$$\text{Ent}(j_t | \rho_t \otimes \kappa) = \int_V \text{Ent}(\hat{\kappa}_t(x, \cdot) | \kappa(x, \cdot)) \rho_t(dx),$$

which is precisely the inner integrand of the dynamic cost in (3.7). Fixing the initial datum  $\rho_0 = \bar{\rho}$  so that the initial term vanishes, (3.1) becomes an action on curves of measures and fluxes with Lagrangian  $L$ , i.e.,

$$\mathcal{I}_{LD}(\rho, j) = \int_0^T L(\rho_t, j_t) dt, \quad L(\rho, j) = \text{Ent}(j | \rho \otimes \kappa), \quad \text{eq:LD-action} \quad (3.8)$$

finite only along pairs  $(\rho, j)$  satisfying the continuity equation

$$\dot{\rho} + \text{div } j = 0, \quad (\text{div } j)(dx) := \int_{y \in V} [j(dx, dy) - j(dy, dx)]. \quad \text{eq:ce-ld} \quad (\text{CE})$$

The presentation here follows [13, App. A], where it is carried out for the Kac process.

**Remark 3.2** (The free energy is a rate function too) rem:static-sanov Lecture 1 justified the relative entropy  $\text{Ent}(\cdot | \pi)$  as driving energy by appealing to Sanov's theorem for  $N \in \mathbb{N}$  i.i.d. samples  $Y^1, \dots, Y^N$  from the invariant measure  $\pi$ . That is the *static* counterpart of (3.1): writing  $\mathbf{S}^N := \frac{1}{N} \sum_{i=1}^N \delta_{Y^i} \in \mathcal{P}(V)$  for the empirical measure of such samples—a measure on *states*, where  $\mathbf{P}^N$  is a measure on *paths*—one has

$$\text{Prob}(\mathbf{S}^N \approx \rho) \asymp e^{-N \text{Ent}(\rho | \pi)}.$$

So *both* ingredients of the gradient structure are large-deviation rate functions of the same particle system—the free energy from its equilibrium statistics, the dissipation from its path statistics. The splitting discussed in §3.2 below is precisely the statement that the dynamic one contains the static one.

## 3.2 Time reversal performs the splitting

sec:reversal Say that a gradient structure  $(\mathcal{F}, \mathcal{R}, \mathcal{R}^*)$  is *induced by* the particle system if the large-deviation action (3.8) is, up to a positive constant, the energy–dissipation functional (EDB) of Lecture 2:

$$\begin{aligned} \mathcal{I}_{LD}(\rho, j) &= c \mathcal{L}(\rho, j; [0, T]), \\ \mathcal{L}(\rho, j; [0, T]) &= \int_0^T [\mathcal{R}(\rho_t, j_t) + \mathcal{D}(\rho_t)] dt + \mathcal{F}(\rho_T) - \mathcal{F}(\rho_0). \end{aligned}$$

Such an identification gives two things at once. Nonnegativity of  $\mathcal{L}$  needs no proof, since  $\mathcal{I}_{LD}$  is a relative entropy; and the null-minimizers of  $\mathcal{L}$ , which is how Lecture 2 defines solutions, are exactly the typical trajectories of the particle system. The gradient structure is then not postulated but *derived*.

It remains to say why a rate functional should split in this way at all. The argument below is formal but complete, and produces the two halves of  $\mathcal{L}$  with  $c = \frac{1}{2}$ .

**The time-reversal functional.** Let  $\mathbf{s}(x, y) := (y, x)$  denote the swap on  $E$  and set

$$j^\dagger := \mathbf{s}_\# j,$$

the flux of the *time-reversed* observation: if  $j$  records mass moving from  $x$  to  $y$ , then  $j^\dagger$  records the same mass moving from  $y$  to  $x$ . The divergence in (CE) antisymmetrizes, so swapping flips its sign,

$$\operatorname{div} j^\dagger = -\operatorname{div} j,$$

and consequently, if  $(\rho, j)$  solves (CE) on  $[0, T]$ , then so does the genuinely reversed curve  $(\hat{\rho}_t, \hat{j}_t) := (\rho_{T-t}, j_{T-t}^\dagger)$ :

$$\partial_t \hat{\rho}_t = -(\partial_t \rho)(T - t) = \operatorname{div} j_{T-t} = -\operatorname{div} \hat{j}_t.$$

Reversal thus maps admissible pairs to admissible pairs, and it is an involution. Since the action (3.8) integrates a Lagrangian evaluated pointwise in time, the substitution  $t \mapsto T - t$  shows that the price the particle system charges for the reversed observation is

$$\mathcal{J}_{LD}(\hat{\rho}, \hat{j}) = \int_0^T L(\rho_{T-t}, j_{T-t}^\dagger) dt = \int_0^T L(\rho_t, j_t^\dagger) dt =: \mathcal{J}_{LD}^\dagger(\rho, j).$$

This is the *time-reversal functional*. It is defined on the same class of pairs as  $\mathcal{J}_{LD}$ , and  $\mathcal{J}_{LD}^{\dagger\dagger} = \mathcal{J}_{LD}$ . Now split  $\mathcal{J}_{LD}$  into its symmetric and antisymmetric parts under  $\dagger$ :

$$\mathcal{J}_{LD} = \underbrace{\frac{1}{2}(\mathcal{J}_{LD} + \mathcal{J}_{LD}^\dagger)}_{\text{symmetric}} + \underbrace{\frac{1}{2}(\mathcal{J}_{LD} - \mathcal{J}_{LD}^\dagger)}_{\text{antisymmetric}}.$$

For a process satisfying detailed balance with respect to  $\pi$ , these two pieces are precisely the two halves of an energy–dissipation balance, i.e., formally,

$$\begin{aligned} \frac{1}{2}(\mathcal{J}_{LD} + \mathcal{J}_{LD}^\dagger) &= \frac{1}{2} \int_0^T [\mathcal{R}(\rho_t, j_t) + \mathcal{D}(\rho_t)] dt, \\ \frac{1}{2}(\mathcal{J}_{LD} - \mathcal{J}_{LD}^\dagger) &= \frac{1}{2}(\mathcal{F}(\rho_T) - \mathcal{F}(\rho_0)). \end{aligned} \tag{3.9} \quad \text{eq:antisym-part}$$

**The computation.** Both identities come out of one change of variables. Take the reference measure of §3.1 to be the invariant one,  $\bar{\rho} = \pi$ , and assume the detailed-balance condition (DBC) with respect to it. The *edge measure*

$$\vartheta := \pi \otimes \kappa, \quad \vartheta(dx dy) = \pi(dx) \kappa(x, dy)$$

on  $E$  is then symmetric,  $\mathbf{s}_\# \vartheta = \vartheta$ , and we write  $\rho_t = u_t \pi$  and  $j_t = w_t \vartheta$ .

The second is legitimate wherever  $L(\rho_t, j_t) < +\infty$ , since then  $j_t \ll \rho_t \otimes \kappa$ , and the expected flux of Step 3 is

$$(\rho \otimes \kappa)(dx dy) = u(x) \vartheta(dx dy),$$

i.e., its density against  $\vartheta$  is the density of  $\rho$  at the *tail* of the edge. Since  $j_t = \rho_t \otimes \hat{\kappa}_t$  with  $\hat{\kappa}_t = r_t \kappa$  by (3.2),  $w_t(x, y) = u_t(x) r_t(x, y)$  is exactly the tilt, and the Lagrangian below is the dynamic cost (3.6) rewritten w.r.t.  $\vartheta$ . So, with  $\Phi(s) = s \log s - s + 1$ , we obtain

$$\begin{aligned} L(\rho, j) &= \text{Ent}(j | \rho \otimes \kappa) = \int_E \Phi\left(\frac{w(x, y)}{u(x)}\right) u(x) \vartheta(dx dy) \\ &= \int_E \left[ w(x, y) \log \frac{w(x, y)}{u(x)} - w(x, y) + u(x) \right] \vartheta(dx dy). \end{aligned}$$

Because  $\vartheta$  is symmetric,  $j^\dagger = \mathbf{s}_\#(w\vartheta) = (w \circ \mathbf{s}) \vartheta$  has density  $w(y, x)$  at  $(x, y)$ . Relabelling  $(x, y) \leftrightarrow (y, x)$  under  $\vartheta$  then gives

$$L(\rho, j^\dagger) = \int_E \left[ w(x, y) \log \frac{w(x, y)}{u(y)} - w(x, y) + u(y) \right] \vartheta(dx dy).$$

The two Lagrangians differ only in whether the observed flux along an edge is priced against the density at the *tail*  $x$  or at the *head*  $y$ . This asymmetry is the whole content of the splitting, and both follow by subtracting and adding.

*Antisymmetric part.* Subtracting, the  $w \log w$  and  $-w$  terms cancel and

$$\frac{1}{2}(L(\rho, j) - L(\rho, j^\dagger)) = \frac{1}{2} \int_E \left[ w(x, y) \log \frac{u(y)}{u(x)} + u(x) - u(y) \right] \vartheta(dx dy) = \frac{1}{2} \langle \nabla \log u, j \rangle,$$

since  $\int_E [u(x) - u(y)] \vartheta = 0$  by the symmetry of  $\vartheta$ . What is left is exactly half the integrand of the chain rule (CR) for the free energy  $\mathcal{F} = \text{Ent}(\cdot | \pi)$ , whose density  $\Phi$  has  $\Phi'(u) = \log u$ . Integrating in time gives (3.9) and shows that the antisymmetric part sees only the endpoints.

*Symmetric part.* Adding instead, and using  $\frac{1}{2}(\log u(x) + \log u(y)) = \log \sqrt{u(x)u(y)}$ ,

$$\frac{1}{2}(L(\rho, j) + L(\rho, j^\dagger)) = \int_E \left[ w(x, y) \log \frac{w(x, y)}{\sqrt{u(x)u(y)}} - w(x, y) + \frac{u(x) + u(y)}{2} \right] \vartheta(dx dy).$$

This is where the geometric mean  $\mathbf{m}(u, v) = \sqrt{uv}$  comes from and is therefore not a modeling choice but the output of symmetrizing. Using the shorthand  $\mathbf{m} := \sqrt{u(x)u(y)}$  and completing the square in  $\sqrt{u}$ ,

$$\frac{1}{2}(L(\rho, j) + L(\rho, j^\dagger)) = \int_E \mathbf{m} \Phi\left(\frac{w}{\mathbf{m}}\right) \vartheta + \frac{1}{2} \int_E \left( \sqrt{u(x)} - \sqrt{u(y)} \right)^2 \vartheta,$$

where we use the fact that  $\mathbf{m} \Phi(w/\mathbf{m}) = w \log(w/\mathbf{m}) - w + \mathbf{m}$  for the first term, and for the second,  $\frac{1}{2}(\sqrt{u(x)} - \sqrt{u(y)})^2 = \frac{1}{2}(u(x) + u(y)) - \mathbf{m}$ . These two terms are the two dissipation terms. The second is half the cosh Fisher information of Example 2.11, whose integrand is  $2(\sqrt{u} - \sqrt{v})^2$ .

# References

- [1] L. Ambrosio, N. Gigli and G. Savaré, *Gradient flows in metric spaces and in the space of probability measures*, second edition, Lectures in Mathematics ETH Zürich, Birkhäuser, Basel, 2008.
- [2] S.-N. Chow, W. Huang, Y. Li and H. Zhou, Fokker–Planck equations for a free energy functional or Markov process on a graph, *Arch. Ration. Mech. Anal.* **203** (2012), no. 3, 969–1008.
- [3] G. Conforti and C. Léonard, Time reversal of Markov processes with jumps under a finite entropy condition, *Stochastic Process. Appl.* **144** (2022), 85–124.
- [4] M. Erbar and J. Maas, Ricci curvature of finite Markov chains via convexity of the entropy, *Arch. Ration. Mech. Anal.* **206** (2012), no. 3, 997–1038.
- [5] N. Fournier and S. Méléard, A microscopic probabilistic description of a locally regulated population and macroscopic approximations, *Ann. Appl. Probab.* **14** (2004), no. 4, 1880–1919.
- [6] J. Hoeksema and O. Tse, Generalized gradient structures for measure-valued population dynamics and their large-population limit, *Calc. Var. Partial Differential Equations* **62** (2023), no. 5, Paper No. 158.
- [7] A. Hraivoronska and O. Tse, Diffusive limit of random walks on tessellations via generalized gradient flows, *SIAM J. Math. Anal.* **55** (2023), no. 4, 2948–2995.
- [8] J. Maas, Gradient flows of the entropy for finite Markov chains, *J. Funct. Anal.* **261** (2011), no. 8, 2250–2292.
- [9] A. Mielke, A gradient structure for reaction–diffusion systems and for energy-drift-diffusion systems, *Nonlinearity* **24** (2011), no. 4, 1329–1346.
- [10] A. Mielke, On evolutionary  $\Gamma$ -convergence for gradient systems, in *Macroscopic and Large Scale Phenomena*, *Lect. Notes Appl. Math. Mech.* **3**, Springer, 2016, 187–249.
- [11] A. Mielke, A. Montefusco and M. A. Peletier, Exploring families of energy-dissipation landscapes via tilting: three types of EDP convergence, *Contin. Mech. Thermodyn.* **33** (2021), no. 3, 611–637.
- [12] A. Mielke, M. A. Peletier and D. R. M. Renger, On the relation between gradient flows and the large-deviation principle, with applications to Markov chains and diffusion, *Potential Anal.* **41** (2014), no. 4, 1293–1327.
- [13] A. Montefusco, U. Sharma and O. Tse, Fourier–Cattaneo equation: stochastic origin, variational formulation, and asymptotic limits, *Commun. Math. Sci.* **24** (2026), no. 1, 1–42.
- [14] M. A. Peletier, R. Rossi, G. Savaré and O. Tse, Jump processes as generalized gradient flows, *Calc. Var. Partial Differential Equations* **61** (2022), no. 1, Paper No. 33, 85 pp.
- [15] E. Sandier and S. Serfaty, Gamma-convergence of gradient flows with applications to Ginzburg–Landau, *Comm. Pure Appl. Math.* **57** (2004), no. 12, 1627–1672.