

Risk-Entropic Flow Matching

Vahid R. Ramezani* Benjamin England†

November 2025

Abstract

Tilted (entropic) risk, obtained by applying a log-exponential transform to a base loss, is a well established tool in statistics and machine learning for emphasizing rare or high loss events while retaining a tractable optimization problem. In this work, our aim is to interpret its structure for Flow Matching(FM). FM learns a velocity field that transports samples from a simple source distribution to data by integrating an ODE. In rectified FM, training pairs are obtained by linearly interpolating between a source sample and a data sample, and a neural velocity field is trained to predict the straight line displacement using a mean squared error loss. This squared loss collapses all velocity targets that reach the same space-time point into a single conditional mean, thereby ignoring higher order conditional information (variance, skewness, multi-modality) that encodes fine geometric structure about the data manifold and minority branches.

We apply the standard risk-sensitive (log-exponential) transform to the conditional FM loss and show that the resulting tilted risk loss is a natural upper-bound on a meaningful conditional entropic FM objective defined at each space-time point. Furthermore, we show that a small order expansion of the gradient of this conditional entropic objective yields two interpretable first order corrections: covariance preconditioning of the FM residual, and a skew tail term that favors asymmetric or rare branches. On synthetic data designed to probe ambiguity and tails, the resulting risk-sensitive loss improves statistical metrics and recovers geometric structure more faithfully than standard rectified FM.

*Email: vahid@engramsciences.com

†Email: ben@engramsciences.com

1 Introduction

Flow matching (FM) has recently emerged as a powerful alternative to score based diffusion models for building continuous time generative models. Most formulations start from a prescribed probability path $(p_t)_{t \in [0,1]}$ and construct a vector field $u_t(x)$ such that the associated probability flow ODE transports a simple source distribution p_0 to a complex target p_1 [1, 2, 3]. A neural network $u_t^\theta(x)$ is trained by minimizing a supervised squared loss between the model field and a microscopic target velocity defined along conditional paths.

The training objective is an L^2 regression problem in which a model vector field $v_t(x)$ is fit to a random microscopic velocity $u_t(X_t | X_1)$ under the path measure. Formally, classical FM solves, at each (t, x) , the L^2 projection problem

$$u_t(x) = \mathbb{E}[u_t(X_t | X_1) | X_t = x],$$

so that the target velocity field is the conditional mean of a richer ensemble of conditional paths. The associated ODE or probability flow construction then serves as a mechanism for aggregating these local regression problems into a global generative model with the correct marginals in the idealized limit of a perfectly fitted vector field.

This is not to say that FM, as a generative framework, is merely a regression estimator. Flow matching has recently been shown to achieve almost minimax optimal convergence rates under strong smoothness and regularity assumptions on the data distribution and the interpolation path [12, 13, 11]. However, these guarantees typically rely on globally smooth, light-tailed densities, globally Lipschitz velocity fields, and strong regularity that are hard to verify in practice and can lead to very loose constants in high dimension. Real-world data manifolds are highly complex and far from globally smooth. Empirical studies of natural-image statistics consistently report sharply peaked, heavy-tailed filter-response distributions [18, 10]. Furthermore, the bounds, naturally, depend on the size of datasets.

In many areas of statistics, control, and finance, it is well understood that one can move beyond purely mean based criteria by replacing the plain expectation of a loss with a *risk-sensitive* or *tilted* risk [19] [6] [8], for example via the log-exponential (entropic) transform as in [5, 15, 4, 16, 7] and recent work on tilted empirical risk in machine learning [6, 9].

This endows the same underlying loss with a different aggregation func-

tional that emphasizes variability, tails, or rare events, and has been studied both for its decision theoretic properties and for its robustness in learning. What is its interpretation for FM? Answering this question is the main contribution of this paper.

In this work, we apply this entropic risk principle to the FM loss, and show that the resulting structure is particularly tractable. At the local level, we introduce a *conditional entropic flow matching objective* obtained by applying a log-exponential transform to the squared FM residual $\|u_t^\theta(x) - U_t(x, z)\|^2$. The gradient of this objective retains a clean FM form: it can be written as the usual FM update with the target velocity replaced by a Gibbs tilted conditional mean $m_t^\lambda(x; \theta)$. A small λ expansion of this tilted mean reveals two interpretable first order corrections to standard FM: a covariance preconditioning term that rescales the FM residual by the conditional covariance of the microscopic velocity, and a skew/tail term that biases updates towards asymmetric or rare branches of the conditional path ensemble. Furthermore, optimizing the conditional entropic objective directly is computationally inconvenient, but it admits the marginal *tilted risk* as an upper-bound.

We begin with some background on rectified Flow Matching following the development given in [3]. Assume $p_0(\xi_0) = \mathcal{N}(\xi_0; x_0, I)$ and $p_1(\xi_1) = \mathcal{N}(\xi_1; x_1, I)$ with fixed means $x_0, x_1 \in \mathbb{R}^d$, and consider a constant velocity field $v_\theta(\xi_t, t) \equiv \theta$. The likelihood of a data point x_1 can be assessed via the instantaneous change of variables formula [2]

$$\log p_1(x_1) = \log p_0(x_0) + \int_0^1 \operatorname{div} v_\theta(x_t, t) dt. \quad (1)$$

Since $\operatorname{div} v_\theta \equiv 0$, the instantaneous change of variables formula (1) reduces to

$$\log p_1(\xi_1) - \log p_0(\xi_0) = 0, \quad (2)$$

on the other hand, by definition:

$$\xi_1 = \xi_0 + \int_0^1 v_\theta(\xi_t, t) dt = \xi_0 + \theta. \quad (3)$$

Using the Gaussian forms of p_0 and p_1 , the identity (3) is equivalent to

$$(\xi_0 - x_0)^\top (\xi_0 - x_0) - (\xi_1 - x_1)^\top (\xi_1 - x_1) \equiv 0 \quad \forall \xi_0, \xi_1. \quad (4)$$

Substituting $\xi_1 = \xi_0 + \theta$ yields, for all $\xi_0 \in \mathbb{R}^d$,

$$0 = (\xi_0 - x_0)^\top (\xi_0 - x_0) - (\xi_0 + \theta - x_1)^\top (\xi_0 + \theta - x_1) \quad (5)$$

$$= (x_1 - x_0 - \theta)^\top (2\xi_0 - x_0 - x_1 + \theta). \quad (6)$$

Since this must hold for every ξ_0 , we must have $x_1 - x_0 - \theta = 0$, i.e.

$$\theta = x_1 - x_0. \quad (7)$$

Consequently, the flow induced by the ODE $\dot{x}_t = v_\theta(x_t, t) \equiv \theta$ is

$$x_t = x_0 + t\theta = (1 - t)x_0 + tx_1, \quad (8)$$

which is precisely the rectified (linear) interpolation between x_0 and x_1 .

In this uni-modal setting, it is natural to model the ground truth velocity as a Gaussian observation

$$x_1 - x_0 \sim \mathcal{N}(v_\theta(x_t, t), I), \quad (9)$$

so that the negative log-likelihood is, up to an additive constant,

$$\mathcal{L}(\theta) = \mathbb{E} \left[\|v_\theta(x_t, t) - (x_1 - x_0)\|_2^2 \right]. \quad (10)$$

Thus, the usual squared loss used in rectified flow matching coincides with maximum likelihood for this Gaussian observation model and is therefore principled in this unimodal case.

The simple setup provides insight into the objective of classic rectified flow matching and offers an alternative way to interpret the flow matching procedure. Instead of two single Gaussian probability density, we now assume that the true source and target distributions are empirical mixtures of Gaussians centered at the data points. For example, for a finite dataset $\mathcal{S} \subset \mathbb{R}^d$ we may write

$$p_0(\xi_0) = \frac{1}{|\mathcal{S}_0|} \sum_{x_0 \in \mathcal{S}_0} \mathcal{N}(\xi_0; x_0, I), \quad p_1(\xi_1) = \frac{1}{|\mathcal{S}_1|} \sum_{x_1 \in \mathcal{S}_1} \mathcal{N}(\xi_1; x_1, I), \quad (11)$$

with $\mathcal{S}_0, \mathcal{S}_1$ the source and target samples respectively.

Let's assume that the velocity vector field $v_\theta(x_t, t)$ at a space-time location (x_t, t) is characterized by a unimodal Gaussian observation model

$$p(v \mid x_t, t) = \mathcal{N}(v; v_\theta(x_t, t), I), \quad (12)$$

with parametric mean $v_\theta(x_t, t)$ and unit covariance. In the rectified flow setting, a pair (x_0, x_1) is connected by the linear trajectory

$$x_t = (1 - t)x_0 + tx_1, \quad (13)$$

so that the associated empirical velocity is simply $x_1 - x_0$, which is independent of t for a given pair. Maximizing the log-likelihood of these empirical velocity observations at (x_t, t) leads to the objective

$$\mathbb{E}_{t, x_0, x_1} [\log p(x_1 - x_0 \mid x_t, t)] \propto -\mathbb{E}_{t, x_0, x_1} \|v_\theta(x_t, t) - (x_1 - x_0)\|_2^2, \quad (14)$$

where the expectation is over $t \sim \mathcal{U}[0, 1]$ and data pairs (x_0, x_1) . Thus, up to an additive constant and a positive scale factor, the negative log-likelihood coincides with the classic rectified flow matching loss:

$$\mathcal{L}_{\text{RF}}(\theta) = \mathbb{E}_{t, x_0, x_1} \|v_\theta(x_t, t) - (x_1 - x_0)\|_2^2. \quad (15)$$

In other words, the standard squared loss used in rectified flow matching can be seen as a principled maximum likelihood estimator for the (Gaussian) observation model on velocities. This interpretation also makes clear that the MSE objective is tailored to the case where the conditional law of the true velocity $x_1 - x_0$ at (x_t, t) is effectively unimodal and Gaussian. In regimes where the conditional distribution is multimodal, heteroscedastic, or skewed, the MSE fit tends to average over distinct plausible directions, yielding a compromise vector field that can under-represent rare but valid branches.

The rest of the paper is organized as follows. First, we briefly review rectified flow matching. Next, we introduce an application of the entropic risk measure [15, 5] to the FM loss. We demonstrate that a first order expansion in the risk coefficient (small λ) corresponds to a modification of the gradient of the FM parameterized vector field. This conditional formulation, while theoretically grounded, does not lend itself directly to efficient sampling for training. Fortunately, a closely related unconditional risk-sensitive loss is easier to sample from and we adopt this practical surrogate in our experiments.

2 Background: Flow Matching and Rectified Paths

Let $(\mathcal{X}, \langle \cdot, \cdot \rangle)$ be Euclidean with norm $\|\cdot\|$. A (time dependent) velocity field $u_\theta^t : \mathcal{X} \rightarrow \mathcal{X}$ induces a path of probability measures $(p_t)_{t \in [0, 1]}$ via the

continuity equation

$$\partial_t p_t(x) + \nabla_x \cdot (p_t(x) u_\theta^t(x)) = 0, \quad (16)$$

with prescribed source p_0 and terminal distribution p_1 matching the data.

Conditional Flow Matching (CFM) specifies a training distribution over space-time points and local velocity targets. We will use z in the rest of this paper for data points used for the construction of conditional paths. In the flow-matching literature, $U_t(x, z)$ is often written as $U_t(x | z)$ and is referred to as a conditional vector field. However, the condition is just a reminder that $U_t(x | z)$ is the vector field that generates $p_t(\cdot | z)$ and not an actual probabilistic condition. The CFM objective is the mean-squared regression, which we denote with MSE throughout:

$$\mathcal{L}_{\text{MSE}}(\theta) = \mathbb{E} \left[\|u_\theta^t(x) - U_t(x, z)\|^2 \right], \quad (17)$$

which trains the parametric field u_θ^t to match the conditional velocity field $U_t(x, z)$ along the path p_t . Rectified Flow Matching is obtained by a particular choice of training pair construction defining a constant velocity $U_t(x, z) = x_1 - x_0$:

$$\mathcal{L}_{\text{FM}}(\theta) = \mathbb{E}_{t, x_0, x_1} \left[\|u_\theta^t(x_t) - (x_1 - x_0)\|^2 \right], \quad (18)$$

In Conditional Flow Matching, the conditional mean velocity $\mu_t(x) = \mathbb{E}[U_t | X_t = x]$ is the transport field that renders the path p_t a solution of the continuity equation. Our risk tilted objective changes only the *local* averaging rule, as shown below, the same gradient form appears with μ_t replaced by a tilted mean.

3 Risk-Sensitive Flow Matching

In this section, we introduce Risk-Sensitive Flow Matching. We show that for small values of risk-sensitive/temperature parameter λ , the risk loss perturbation amounts to a local higher order perturbation of the velocity field.

Define the centered residual

$$\varepsilon_t := U_t - \mu_t(X_t), \quad \mathbb{E}[\varepsilon_t | X_t] = 0. \quad (19)$$

where

$$\mu_t(x) := \mathbb{E}[U_t | X_t = x], \quad (20)$$

Gradient of the mean squared objective. Consider the squared loss

$$\mathcal{L}_{\text{MSE}}(\theta) = \mathbb{E}[\|u_\theta^t(X_t) - U_t\|^2], \quad (21)$$

where $u_\theta^t : \mathcal{X} \rightarrow \mathbb{R}^d$ is differentiable in θ and integrable so that differentiation under the integral is valid. Writing $u_\theta := u_\theta^t$ for brevity, the chain rule gives

$$\nabla_\theta \|u_\theta(X_t) - U_t\|^2 = 2(u_\theta(X_t) - U_t)^\top \nabla_\theta u_\theta(X_t).$$

Taking expectations and applying the tower property,

$$\begin{aligned} \nabla_\theta \mathcal{L}_{\text{MSE}}(\theta) &= 2 \mathbb{E}[(u_\theta(X_t) - U_t)^\top \nabla_\theta u_\theta(X_t)] \\ &= 2 \mathbb{E}\left[\mathbb{E}[(u_\theta(X_t) - U_t)^\top \nabla_\theta u_\theta(X_t) \mid X_t, t]\right]. \end{aligned}$$

Conditioning on $(X_t = x)$ and using $U_t = \mu_t(x) + \varepsilon_t$ with $\mathbb{E}[\varepsilon_t \mid X_t = x] = 0$,

$$\mathbb{E}[(u_\theta(x) - U_t)^\top \nabla_\theta u_\theta(x) \mid X_t = x] = (u_\theta(x) - \mu_t(x))^\top \nabla_\theta u_\theta(x).$$

Therefore,

$$\nabla_\theta \mathcal{L}_{\text{MSE}}(\theta) = 2 \mathbb{E}[(u_\theta^t(X_t) - \mu_t(X_t))^\top \nabla_\theta u_\theta^t(X_t)]. \quad (22)$$

Stationarity and mean collapse. In the unconstrained (nonparametric) function class, the global minimizer of $\mathcal{L}_{\text{MSE}}$ is

$$u^*(x, t) = \mu_t(x) \quad \text{a.s.}$$

3.1 Risk-Sensitive Objective

The log-exponential (entropic) transformation of a nonnegative loss $C(\theta)$ is defined as:

$$\mathcal{L}_\lambda(\theta) := \frac{1}{\lambda} \log \mathbb{E}[\exp(\lambda C(\theta))], \quad \lambda > 0. \quad (23)$$

It smoothly interpolates between the standard mean loss and a soft maximum. This follows from the convexity and smoothness of the log-sum-exp function and its role as the cumulant generating transform.

(23) admits an information theoretic representation [16, pp. 174]. Let P be the data distribution and $Q \ll P$ any reweighting. Then

$$\frac{1}{\lambda} \log \mathbb{E}_P[e^{\lambda C}] = \sup_{Q \ll P} \left\{ \mathbb{E}_Q[C] - \frac{1}{\lambda} D_{\text{KL}}(Q \| P) \right\}. \quad (24)$$

Thus minimizing $\mathcal{L}_\lambda$ is equivalent to minimizing the worst case expected loss over all KL bounded reweightings of P . As λ increases, the optimizer Q^* tilts toward high loss samples; as $\lambda \downarrow 0$, the KL penalty dominates and $Q^* \rightarrow P$, recovering the usual mean loss. In risk theory, $\mathcal{L}_\lambda$ is the entropic risk measure (exponential utility certainty equivalent), adding variance and tail sensitive emphasis as λ grows.

In our setting:

$$\mathcal{L}_\lambda(\theta) = \frac{1}{\lambda} \log \mathbb{E} \left[\exp(\lambda \|u_\theta^t(x) - U_t(x, z)\|^2) \right].$$

We will show that this loss is an upper bound for a conditional entropic objective in the next section.

3.2 Conditional (Entropic/Gibbs) Objective

Observe that we can write the Conditional Flow Matching loss by reversing the order of conditioning:

$$\begin{aligned} L_{\text{MSE}}(\theta) &= \mathbb{E}_{t,z} \left[\mathbb{E}_{x|t,z} \|u_\theta^t(x) - U_t(x, z)\|^2 \right] \\ &= \mathbb{E}_{t,x,z} [\|u_\theta^t(x) - U_t(x, z)\|^2] \\ &= \mathbb{E}_{x,t} \left[\mathbb{E}_{z|x,t} \|u_\theta^t(x) - U_t(x, z)\|^2 \right]. \end{aligned} \tag{25}$$

Starting from the last (conditional) form, we replace the inner conditional mean over $z | x$ by a log-exp (entropic) aggregation and for $\lambda > 0$, define the conditional risk tilted loss [4] [5]

$$L_\lambda(\theta) := \mathbb{E}_{x,t} \left[\frac{1}{\lambda} \log \mathbb{E}_{z|x,t} \exp(\lambda \|u_\theta^t(x) - U_t(x, z)\|^2) \right], \tag{26}$$

The inner log-exponential is the entropic (Gibbs) risk of the squared error at (x, t) ; it reduces to the mean as $\lambda \rightarrow 0$. This form gives us better insight as to how the higher order moments of the microscopic velocity field manifest themselves. Next, we will give the main results. The detailed proofs are given in Appendix B.

Tilted weights and gradient. Define the conditional Gibbs weights

$$\beta_\lambda(z | x, t; \theta) := \frac{\exp(\lambda \|u_\theta^t(x) - U_t(x, z)\|^2)}{\mathbb{E}_{z|x,t} \exp(\lambda \|u_\theta^t(x) - U_t(x, z)\|^2)}, \quad \mathbb{E}_{z|x,t}[\beta_\lambda] = 1,$$

and the resulting *tilted conditional mean*

$$m_\lambda^t(x; \theta) := \mathbb{E}_{z|x,t} [\beta_\lambda(z | x, t; \theta) U_t(x, z)]. \quad (27)$$

Differentiating under the integral sign yields

$$\nabla_\theta L_\lambda(\theta) = 2 \mathbb{E}_{x,t} \left[(u_\theta^t(x) - m_\lambda^t(x; \theta))^\top \nabla_\theta u_\theta^t(x) \right]. \quad (28)$$

Comparing (22) and (28), the only change from MSE is the replacement

$$\mu_t(x) \quad \longmapsto \quad m_\lambda^t(x; \theta)$$

in the local averaging rule at each (x, t) .

Moment expansion of the tilted mean. Write

$$U_t(x, z) = \mu_t(x) + \varepsilon, \quad \mathbb{E}[\varepsilon | x, t] = 0,$$

and define conditional moments

$$\Sigma_t(x) = \mathbb{E}[\varepsilon \varepsilon^\top | x, t], \quad S_t(x) = \mathbb{E}[\varepsilon \|\varepsilon\|^2 | x, t].$$

Let the residual error be $m_t(x) = u_\theta^t(x) - \mu_t(x)$, Theorem 1 proven in Appendix B shows

$$m_\lambda^t(x; \theta) = \mu_t(x) + \lambda S_t(x) - 2\lambda \Sigma_t(x) m_t(x) + O(\lambda^2). \quad (29)$$

from (28) and (22),

$$\nabla_\theta L_\lambda - \nabla_\theta \mathcal{L}_{\text{MSE}} = 2 \mathbb{E}_{x,t} \left[\left(-\lambda S_t(x) + 2\lambda \Sigma_t(x) m_t(x) \right)^\top \nabla_\theta u_\theta^t(x) \right] + O(\lambda^2). \quad (30)$$

The covariance term $2\lambda \Sigma_t(x) m_t(x)$ acts as a *covariance preconditioner*, scaling updates along directions of large conditional variance and allocating more learning capacity to ambiguous or hard axes. The skew term $\lambda S_t(x)$ acts as a *skew bias*, tilting toward minority or tail branches when the conditional distribution is asymmetric. We can think of this perturbation of the gradient as a form of self-guidance, where, in place of conditioning to modify the underlying probability, we are modifying the probability measure to account for higher order moments of the distribution.

Risk-Sensitive Objective as an upper bound We now show that risk-sensitive objective upper bounds conditional entropic objective. Using tower property of expectations we can write

$$\mathcal{L}_\lambda(\theta) = \frac{1}{\lambda} \log \mathbb{E} \left[\exp(\lambda \|u_\theta^t(x) - U_t(x, z)\|^2) \right].$$

$$\mathcal{L}_\lambda(\theta) = \frac{1}{\lambda} \log \mathbb{E}_{x,t} \left[\mathbb{E}_{z|x,t} \left[\exp(\lambda \|u_\theta^t(x) - U_t(x, z)\|^2) \right] \right].$$

Compared to (26), local higher order moments with respect to z are averaged over (x, t) before they are flattened by the logarithm. However, on average, the objective still relies on the minimization of local higher order structure represented by the inner expectation; the only difference is that the logarithm is applied after this averaging rather than before it. We now make this precise.

Define

$$\phi_\lambda(x, t; \theta) := \frac{1}{\lambda} \log \mathbb{E}_{z|x,t} \exp(\lambda \|u_\theta^t(x) - U_t(x, z)\|^2),$$

so that the conditional objective is

$$L_\lambda(\theta) = \mathbb{E}_{x,t} [\phi_\lambda(x, t; \theta)],$$

Since log is concave, Jensen's inequality for concave functions [17] gives

$$L_\lambda(\theta) = \mathbb{E}_{x,t} [\phi_\lambda(x, t; \theta)] \leq \mathcal{L}_\lambda(\theta),$$

for all $\lambda > 0$. Thus, $\mathcal{L}_\lambda$ is a pointwise (in parameters) upper-bound on the conditional entropic objective L_λ , and in what follows we simply use $\mathcal{L}_\lambda$ as a tractable surrogate, in the same spirit as such bound surrogates are commonly optimized in machine learning.

4 Related Work

To our knowledge, this is the first work to characterize entropic risk in FM setting. We adopt a tilted (log-exponential) risk on the standard FM loss, building on a long line of work on entropic and risk-sensitive objectives in control and mathematical finance [5, 15, 4, 16]. In machine learning, risk-sensitive learning has been studied in a general way by Lee et al. [19], who

analyze learning bounds for objectives defined via optimized certainty equivalents, including the entropic risk as a special case; Li et al. [6] introduce *Tilted Empirical Risk Minimization*, applying exponential tilting directly to the empirical loss distribution to emphasize high or low loss samples; Curi et al. [20] develop an adaptive sampling algorithm that improves performance on a given fraction of most difficult examples. Wu et al. [21] reinterpret contrastive learning objectives such as InfoNCE as distributionally robust optimization, with the temperature parameter acting as a Lagrange multiplier that controls the size of an adversarial distribution set.

A line of recent work on rectified flows analyzes the ensemble of transport paths induced by random source-target couplings and highlights that naive ℓ_2 averaging of velocity targets can be harmful. Variational Rectified Flow Matching [3] makes this critique explicit and introduces latent variables to represent multi-modal velocity fields at each (x_t, t) , allowing distinct branches to cross while remaining consistent with the same marginals. [14] introduces a hierarchical system of ordinary differential equations enabling a more accurate modeling of the underlying velocity distribution. Our work is complementary. Combining such mode aware parameterizations with our risk-sensitive objective that shifts the statistics is an interesting direction for future work.

5 Experiments on Synthetic Data

Our early numerical analysis on synthetic data has been encouraging, we share one of them below, more details can be found in the appendix. For ground truth, we replicated the 2D experiment given in [3]

Setup. We consider a rectified flow matching task on a synthetic two ring Gaussian mixture: sources x_0 are sampled on an inner ring and targets x_1 on an outer ring, each with $K=6$ equally spaced angular lobes and small radial noise. Training uses pairs (x_0, x_1) , a uniform $t \sim \text{Unif}[0, 1]$, rectified states $x_t = (1 - t)x_0 + tx_1$, and the supervised velocity $u_t = x_1 - x_0$. Thus, the ground truth transport consists of straight line chords between paired points on the two rings, as visualized in Figure 1. We compare standard FM (MSE) to the *marginal* entropic FM objective with a scheduled risk coefficient $\lambda(t_{\text{step}})$ that linearly ramps from 0 to λ_{max} over 500 steps and is then held fixed. An MLP with sinusoidal embeddings predicts $v_\theta(x, t)$; all

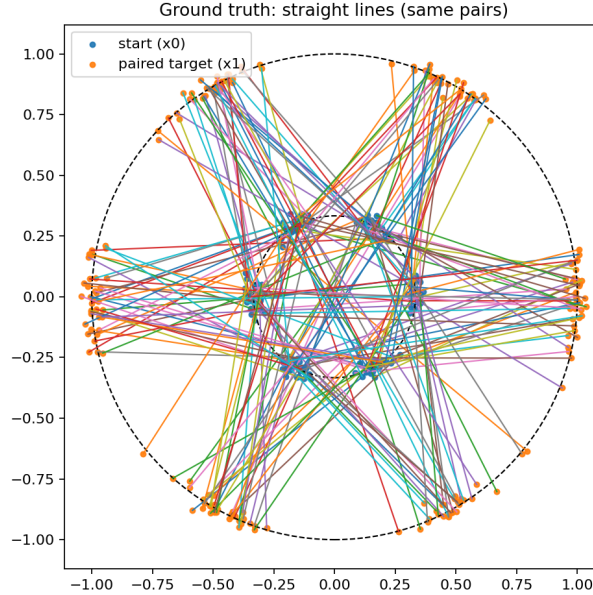


Figure 1: Ground truth straight line transport paths in the synthetic 2D ring experiment. Blue dots mark source points x_0 and orange dots their paired targets x_1 on the ring; each colored segment shows the straight line trajectory $x_t = (1 - t)x_0 + tx_1$. The outer dashed circle indicates the target ring.

runs use identical seeds, optimizer (AdamW), batch size, and step budget.

Metrics. We report four diagnostics: (i) RMSE_σ between the model-recovered lobe spread and the ground truth spread (lower is better); (ii) *gap violation rate*—the fraction of generated angles whose nearest-lobe residual exceeds a calibrated threshold (lower is better); (iii) the mean 1D Wasserstein-1 distance between the empirical distributions of absolute angular residuals for the ground truth and the model, $W_1(|r|)$ (lower is better); and (iv) the model’s estimated spread σ_{model} (for context).

Results (scheduled risk sweep). Across a sweep of $\lambda_{\text{max}} \in [0.01, 0.40]$, the entropic objective consistently improves *variance matching* and *mode concentration* while largely preserving mass away from inter-lobe gaps. The plain FM baseline (corresponding to $\lambda_{\text{max}} = 0$) attains

$$\text{RMSE}_\sigma^{\text{base}} \approx 0.0149, \quad W_1^{\text{base}}(|r|) \approx 0.0121, \quad \text{Gap}^{\text{base}} \approx 0.0297.$$

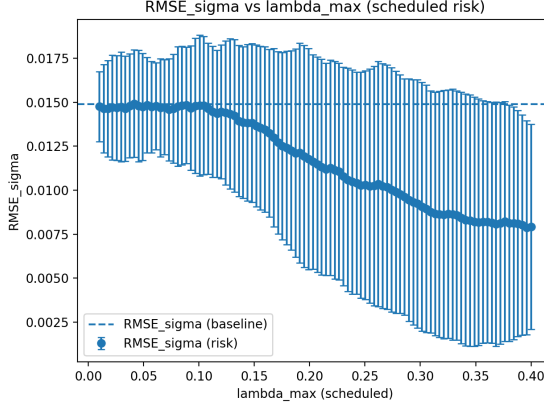


Figure 2: RMSE_σ vs. $\lambda_{\max}$ (FM baseline (dashed line) vs. scheduled FM-RISK. Lower is better).

As shown in Figure 2, RMSE_σ decreases monotonically as $\lambda_{\max}$ increases. Figure 3 summarizes the corresponding relative improvement. For a representative mid-range setting $\lambda_{\max} \approx 0.25$ we already obtain a reduction of about 31% in RMSE_σ relative to the baseline, while the Wasserstein distance Figure 5 improves by roughly 6% and the gap violation rate Figure 4 drops by about 8.5%.

At the upper end of the sweep ($\lambda_{\max} \approx 0.40$) the model reaches $\text{RMSE}_\sigma^{\text{risk}} \approx 0.0079$, corresponding to a relative improvement of roughly 47% over FM. In this more aggressive regime the gap rate remains slightly better than baseline, whereas $W_1(|r|)$ drifts back to around the FM value and becomes marginally worse (about 3% larger), reflecting a trade-off in which the risk term prioritizes shrinking the core angular spread.

Figure 4 reports the gap violation rate as a function of $\lambda_{\max}$. For small coefficients, the behavior is close to FM, but for moderate values the violation rate clearly drops below the baseline, indicating better respect of the underlying angular structure; for very large coefficients the gap rate slightly increases again but remains comparable to or better than FM.

In Figure 5 we plot the mean Wasserstein distance $W_1(|r|)$ between the residual distribution and its reference. The curve exhibits a broad U-shape: moderate risk coefficients reduce $W_1(|r|)$ relative to FM, while extremely large values tilt the objective enough that the distance grows back toward the baseline.

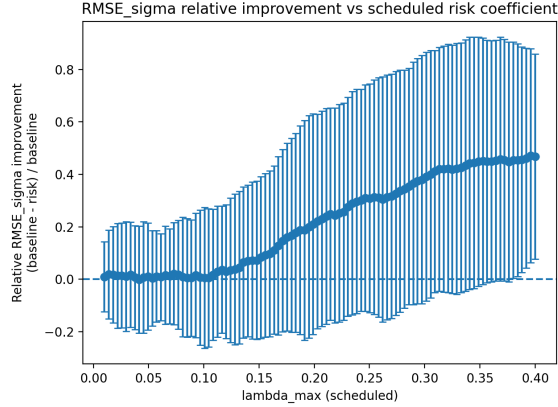


Figure 3: Relative RMSE_σ improvement $(\text{RMSE}_\sigma^{\text{base}} - \text{RMSE}_\sigma^{\text{risk}})/\text{RMSE}_\sigma^{\text{base}}$ vs. λ_{max} (FM baseline (dashed line) vs. scheduled FM-RISK. Higher is better)

Joint improvement regime. Across the sweep there is a broad regime of *moderate* risk coefficients in which all three key angular metrics improve simultaneously relative to the FM baseline. Empirically, for λ_{max} in the range $\lambda_{\text{max}} \approx 0.2\text{--}0.3$, we observe concurrent reductions in RMSE_σ , calibrated gap-violation rate, and $W_1(|r|)$: the angular spread error typically decreases by roughly 30–40%, the gap rate drops by about 5–9%, and the Wasserstein distance on absolute residuals improves by around 5–6% relative to baseline. Outside this band the risk-sensitive objective still yields gains in RMSE_σ , but the improvements in gap rate and $W_1(|r|)$ become weaker or trade off against each other, indicating that moderate risk levels are best aligned with jointly sharpening lobe concentration, reducing gap mass, and matching the residual distribution.

On this simple, "smooth" two-ring benchmark, standard flow matching already achieves, as expected, fairly small angular and gap errors, so the absolute improvements from risk-sensitive training are modest in magnitude, even though the relative reductions in spread error and gap rate are substantial.

6 Limitations and Future Work

Our method introduces an entropic (log-exp) regularizer for flow matching with a temperature $\lambda > 0$. While the first-order analysis clarifies that the in-

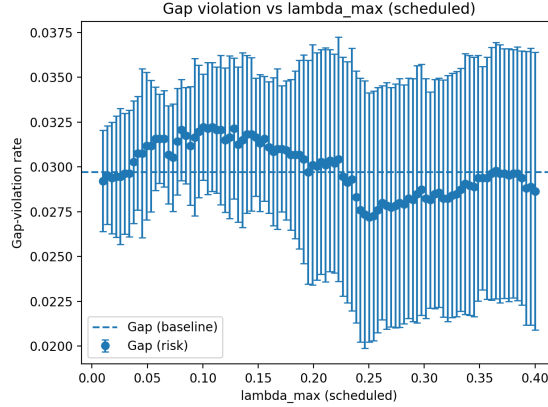


Figure 4: Gap violation rate vs. $\lambda_{\max}$ (FM baseline (dashed line) vs. scheduled FM-RISK. Lower is better).

tended regime is small λ , we do not yet provide a principled rule for choosing or adapting λ .

Hyperparameter and scheduling. We currently set λ by a coarse sweep and found that scheduling it to ramp up later in training helps as early iterations are noise-dominated, so higher order (beyond the conditional mean) information is unreliable. Studying the scheduling of the λ hyperparameter remains an interesting problem.

Scope of experiments. Experiments here are on controlled synthetic families. Broader validation on image/audio/text domains, ablations across network sizes and optimizers, and comparisons to hard-example mining and curriculum schedules are needed to assess robustness and scalability.

7 Conclusion

Tilted (log-exponential) risk on a loss is by now a standard tool in statistics and machine learning, used to emphasize rare or high loss events while preserving a clean optimization structure. In this work, we apply this well established construction directly to flow matching. Classical flow matching can be viewed as solving, at each (t, x) , the L^2 projection problem so that

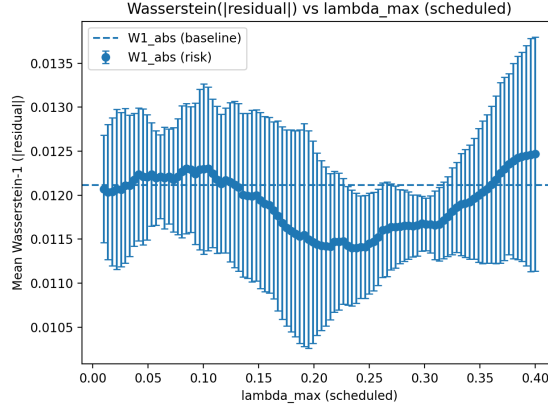


Figure 5: Mean Wasserstein–1 distance on absolute angular residuals $W_1(|r|)$ vs. $\lambda_{\max}$ (FM baseline (dashed line) vs. scheduled FM-RISK. Lower is better).

the target velocity field is the conditional mean of a richer microscopic velocity. Our entropic objective takes this characterization as its starting point and replaces the plain mean square aggregation by a Gibbs weighted one: we apply the tilted (log-exponential) risk to the FM loss and show that, in the FM setting, the resulting global objective is a natural upper-bound on a conditional entropic FM functional defined at each space-time point. This yields a one parameter family of risk entropic flow matching. At small values of this coefficient or temperature λ the objective behaves like standard FM plus a correction, offering robustness and sharper control of higher order statistics; at higher values it might act as an explicit innovation or exploration mechanism, gradually reweighting tails and rare modes and thereby changing the effective measure on the data manifold in a principled way.

References

- [1] Y. Lipman, R. T. Q. Chen, H. Ben Hamu, M. Nickel, and M. Le. Flow matching for generative modeling. In *International Conference on Learning Representations (ICLR)*, 2023. arXiv:2210.02747, 2022.
- [2] Y. Lipman, R. T. Q. Chen, H. Ben Hamu, M. Nickel, and M. Le. *Flow Matching for Generative Modeling*. Advances in Neural Information Processing Systems (NeurIPS), 2023.

- [3] P. Guo and A. G. Schwing. *Variational Rectified Flow Matching*. arXiv preprint arXiv:2502.09616, 2025.
- [4] K. Detlefsen and G. Scandolo. Conditional and dynamic convex risk measures. *Finance and Stochastics*, 9(4):539–561, 2005.
- [5] R. A. Howard and J. E. Matheson. Risk-sensitive Markov decision processes. *Management Science*, 18(7):356–369, 1972.
- [6] T. Li, A. Beirami, M. Sanjabi, and V. Smith. Tilted empirical risk minimization. arXiv preprint arXiv:2007.01162, 2021.
- [7] L. Leqi, A. S. Ross, B. Bartlett, and M. I. Jordan. Supervised learning with general risk functionals. In *International Conference on Machine Learning (ICML)*, 2022.
- [8] T. Sypherd, A. Orlitsky, and V. Kostina. A tunable loss function for robust classification. *IEEE Transactions on Information Theory*, 68(10):6510–6530, 2022.
- [9] G. Aminian, A. R. Asadi, T. Li, A. Beirami, G. Reinert, and S. N. Cohen. Generalization and robustness of the tilted empirical risk. In *International Conference on Machine Learning (ICML)*, volume 267 of *Proceedings of Machine Learning Research*, 2025.
- [10] J. Huang and D. Mumford. Statistics of natural images and models. In *Proc. IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 541–547, 1999.
- [11] K. Fukumizu, T. Suzuki, N. Isobe, K. Oko, and M. Koyama. Flow matching achieves almost minimax optimal convergence. *arXiv preprint arXiv:2405.20879*, 2024.
- [12] M. S. Albergo and E. Vanden-Eijnden. Building normalizing flows with stochastic interpolants. In *International Conference on Learning Representations (ICLR)*, 2023.
- [13] J. Benton, G. Deligiannidis, and A. Doucet. Error bounds for flow matching methods. In *Transactions on Machine Learning Research (TMLR)*, 2024.

- [14] Y. Zhang, Y. Yan, A. Schwing, and Z. Zhao. Towards hierarchical rectified flow. In *Proceedings of the Thirteenth International Conference on Learning Representations (ICLR)*, 2025.
- [15] V. R. Ramezani and S. I. Marcus. Estimation of hidden Markov models: Risk-sensitive filter banks and qualitative analysis of their sample paths. *IEEE Transactions on Automatic Control*, 47(12):1999–2009, 2002.
- [16] H. Follmer and A. Schied. *Stochastic Finance: An Introduction in Discrete Time*. de Gruyter, 2nd edition, 2004.
- [17] F. M. Dekking, C. Kraaikamp, H. P. Lopuhaa, and L. E. Meester. *A Modern Introduction to Probability and Statistics: Understanding Why and How*. Springer, 2005.
- [18] D. L. Ruderman and W. Bialek. Statistics of natural images: Scaling in the woods. In *Advances in Neural Information Processing Systems (NIPS 6)*, pages 551–558. Morgan Kaufmann, 1994.
- [19] J. Lee, S. Park, and J. Shin. Learning bounds for risk-sensitive learning. In *Advances in Neural Information Processing Systems (NeurIPS)*, 33, 2020.
- [20] S. Curi, K. Y. Levy, S. Jegelka, and A. Krause. Adaptive sampling for stochastic risk-averse learning. In *Advances in Neural Information Processing Systems (NeurIPS)*, 33, 2020.
- [21] J. Wu, J. Chen, J. Wu, W. Shi, X. Wang, and X. He. Understanding contrastive learning via distributionally robust optimization. In *Advances in Neural Information Processing Systems (NeurIPS)*, 36, 2023.

8 Appendix A: Synthetic 2D ring experiment

Distributions. We consider a pair of 2D distributions supported on concentric circles with six modes each. The source distribution p_0 is defined as an equally weighted mixture of $K = 6$ Gaussian components with means

$$m_k^{(0)} = r_{\text{src}} \begin{bmatrix} \cos(2\pi k/K) \\ \sin(2\pi k/K) \end{bmatrix}, \quad k \in \{0, \dots, K-1\},$$

where $r_{\text{src}} = 1/3$. To sample from p_0 , we first draw a component index k uniformly from $\{0, \dots, K-1\}$, then add independent angular and radial perturbations. Concretely, if $\alpha_k = 2\pi k/K$ denotes the base angle, we sample

$$\tilde{\alpha} \sim \mathcal{N}(\alpha_k, \sigma_{\text{ang}}^2), \quad \tilde{r} \sim \mathcal{N}(r_{\text{src}}, \sigma_{\text{rad}}^2),$$

with $\sigma_{\text{ang}} = 0.12$ and $\sigma_{\text{rad}} = 0.02$, and set

$$x_0 = \tilde{r} \begin{bmatrix} \cos \tilde{\alpha} \\ \sin \tilde{\alpha} \end{bmatrix}.$$

The target distribution p_1 is constructed analogously, but with radius $r_{\text{tgt}} = 1$, i.e.,

$$m_k^{(1)} = r_{\text{tgt}} \begin{bmatrix} \cos(2\pi k/K) \\ \sin(2\pi k/K) \end{bmatrix}.$$

Angular and radial noise use the same standard deviations $\sigma_{\text{ang}} = 0.12$ and $\sigma_{\text{rad}} = 0.02$, so the source and target differ only by overall scale.

Training tuples and baseline FM. At each optimization step we draw a mini batch of $B = 1000$ i.i.d. pairs (x_0, x_1) , where $x_0 \sim p_0$ and $x_1 \sim p_1$ are sampled independently as described above, together with times $t \sim \text{Unif}[0, 1]$. We form rectified paths

$$x_t = (1-t)x_0 + tx_1, \quad u = x_1 - x_0,$$

and train a neural velocity field $v_\theta(x_t, t)$ using the standard rectified flow matching loss

$$\mathcal{L}_{\text{FM}}(\theta) = \mathbb{E} [\|v_\theta(x_t, t) - u\|_2^2].$$

The baseline FM model is trained for 20,000 iterations with batch size $B = 1000$, AdamW optimizer, learning rate 3×10^{-4} , and default weight decay.

Risk-sensitive objective. For the risk-sensitive variant we keep the same data construction, architecture, optimizer, and training schedule, changing only the loss. Let

$$\ell_b(\theta) = \|v_\theta(x_t^{(b)}, t^{(b)}) - u^{(b)}\|_2^2$$

denote the per-sample squared error for batch index $b \in \{1, \dots, B\}$. For a risk coefficient $\lambda > 0$, the empirical log-exponential loss is

$$\hat{\mathcal{L}}_\lambda(\theta) = \frac{1}{\lambda} \log \left(\frac{1}{B} \sum_{b=1}^B \exp(\lambda \ell_b(\theta)) \right),$$

which implements the entropic risk $\frac{1}{\lambda} \log \mathbb{E}[\exp(\lambda \ell)]$ via a numerically stable log-mean-exp over the mini-batch. As $\lambda \rightarrow 0$ this reduces to the MSE loss.

We use a simple schedule λ_t that ramps the risk coefficient from 0 to a prescribed $\lambda_{\max}$ over the first 500 training iterations and then keeps it fixed:

$$\lambda_t = \begin{cases} 0, & t < 0, \\ \frac{t}{500} \lambda_{\max}, & 0 \leq t < 500, \\ \lambda_{\max}, & t \geq 500. \end{cases}$$

For each value of $\lambda_{\max}$ in a sweep we retrain the model from scratch with random uniform initialization and data seeds to enable paired comparisons against the baseline FM model.

Per mode angular spread metric. To quantify how well the learned flow preserves the multi-modal structure of the target distribution, we measure the angular concentration of mass around each mode on the ring. Let K denote the number of lobes (here $K = 6$), and define the ideal mode centers on the unit circle by

$$\mu_k = \frac{2\pi k}{K}, \quad k \in \{0, \dots, K-1\}.$$

Given a 2D-point $x = (x_1, x_2)^\top$, we first compute its polar angle

$$\theta(x) = \text{atan2}(x_2, x_1) \in [0, 2\pi).$$

We then assign x to its nearest lobe and work in *per-lobe coordinates* by subtracting the closest mode center:

$$k^*(x) = \arg \min_k |\text{wrap}(\theta(x) - \mu_k)|, \quad \delta\theta(x) = \text{wrap}(\theta(x) - \mu_{k^*(x)}),$$

where $\text{wrap}(\cdot)$ maps angles into $(-\pi, \pi]$ by adding or subtracting integer multiples of 2π . Intuitively, $\delta\theta(x)$ is the angular residual of x relative to the center of its nearest mode, so that all K lobes are folded onto a single canonical lobe.

We apply this construction to both the ground-truth samples x_1 and the model generated samples $\hat{x}_1$, obtaining residuals $\delta\theta_{\text{true}}$ and $\delta\theta_{\text{model}}$, respectively. For each K group we then compute the *circular* standard deviation of these residuals,

$$\sigma_{\text{true}} = \text{CircStd}(\delta\theta_{\text{true}}), \quad \sigma_{\text{model}} = \text{CircStd}(\delta\theta_{\text{model}}),$$

using the resultant-length-based definition of circular standard deviation. This measures the angular spread of mass around the nearest mode in a way that is consistent on the circle.

8.1 Metrics and Summary Statistics

Gap violation rate (calibrated). Let K be the number of angular lobes and let the lobe centers be $\mu_j = 2\pi j/K$, $j = 0, \dots, K-1$. For a predicted angle $\theta \in [0, 2\pi)$ define the wrapped distance to the nearest lobe

$$d_K(\theta) = \min_j |\text{wrap}(\theta - \mu_j)|, \quad \text{wrap}(\phi) \in (-\pi, \pi].$$

For each K we first *calibrate* a threshold τ_K from the ground-truth residuals by choosing a high quantile (e.g. the 99th percentile) of $|r|$ for that K . A model prediction with angle θ is a *gap violation* if $d_K(\theta) > \tau_K$; samples beyond this calibrated angular distance lie in inter-lobe regions that are rarely visited by the true distribution. Given a set of predictions $\{\theta_i\}_{i=1}^N$, the overall rate is

$$\text{GapRate} = \frac{1}{N} \sum_{i=1}^N \mathbb{I}[d_K(\theta_i) > \tau_{K(i)}],$$

where $K(i)$ denotes the lobe multiplicity associated with θ_i (here always $K = 6$). Lower is better, as it indicates less mass in gap regions, and the use of τ_K ensures that the notion of a “gap” is calibrated to the ground-truth spread.

Wasserstein distances on angular residuals. For a given K , define nearest-lobe angular residuals for ground truth and model by

$$r_i^{\text{true}} = \text{wrap}(\theta_i^{\text{true}} - \mu_{j^*(\theta_i^{\text{true}})}), \quad r_i^{\text{model}} = \text{wrap}(\theta_i^{\text{model}} - \mu_{j^*(\theta_i^{\text{model}})}).$$

We consider the empirical distributions of $|r|$ and of signed r and compute their one-dimensional Wasserstein-1 distances:

$$\begin{aligned} W_1^{(K)}(|r|) &= W_1(\{|r_i^{\text{true}}|\}_i, \{|r_i^{\text{model}}|\}_i), \\ W_1^{(K)}(r) &= W_1(\{r_i^{\text{true}}\}_i, \{r_i^{\text{model}}\}_i), \end{aligned}$$

where W_1 denotes the usual earth-mover distance on $\mathbb{R}$. We then average these per- K distances with weights proportional to the number of samples in each bin:

$$W_1(|r|) = \sum_K w_K W_1^{(K)}(|r|), \quad W_1(r) = \sum_K w_K W_1^{(K)}(r),$$

with $\sum_K w_K = 1$. Lower values indicate that the model reproduces the shape of the residual distribution more closely, both in magnitude and in signed orientation.

Spread error RMSE_σ . Let $\sigma_{\text{true}}(K)$ be the circular standard deviation of the nearest-lobe residuals for ground truth and $\sigma_{\text{model}}(K)$ for the model. With weights w_K (again proportional to per-bin counts),

$$\text{RMSE}_\sigma = \left(\sum_K w_K (\sigma_{\text{model}}(K) - \sigma_{\text{true}}(K))^2 \right)^{1/2}.$$

This summarizes how well the model matches the angular spread of each lobe; lower is better.

Radial error metrics. To ensure that improvements in angular statistics do not come at the expense of radial transport, we also track radial mean-squared and mean-absolute errors. Let $r_i^{\text{true}} = \|x_1^{(i)}\|_2$ and $r_i^{\text{model}} = \|\hat{x}_1^{(i)}\|_2$ denote the radii of ground truth and model samples, respectively. We define

$$\text{RadialMSE} = \frac{1}{N} \sum_{i=1}^N (r_i^{\text{model}} - r_i^{\text{true}})^2,$$

$$\text{RadialMAE} = \frac{1}{N} \sum_{i=1}^N |r_i^{\text{model}} - r_i^{\text{true}}|.$$

These metrics are reported alongside the angular diagnostics but play a supporting role: in our experiments, risk-sensitive training improves angular behaviour while leaving radial errors essentially unchanged or slightly better.

Appendix B: Tilted log-exponential gradients

In this appendix we give a more detailed derivation of the gradient of the conditional log-exponential (entropic) FM objective and its first-order expansion in the tilt parameter λ . Throughout, $\lambda > 0$ is a scalar parameter. The case $\lambda = 0$ recovers the usual FM/MSE objective; small positive values of λ define a deformation of that objective. But first, we need a lemma:

Lemma 1 (First-order ratio expansion). *Let*

$$N(\lambda) = N_0 + \lambda N_1 + O(\lambda^2), \quad D(\lambda) = D_0 + \lambda D_1 + O(\lambda^2),$$

with $D_0 \neq 0$. Then

$$\frac{N(\lambda)}{D(\lambda)} = \frac{N_0}{D_0} + \lambda \frac{N_1 D_0 - N_0 D_1}{D_0^2} + O(\lambda^2).$$

Proof. Starting from

$$D(\lambda) = D_0 + \lambda D_1 + O(\lambda^2),$$

factor out D_0 :

$$D(\lambda) = D_0 \left(1 + \lambda \frac{D_1}{D_0} + O(\lambda^2) \right).$$

Set

$$x(\lambda) := \lambda \frac{D_1}{D_0} + O(\lambda^2),$$

so $x(\lambda) = O(\lambda)$ as $\lambda \rightarrow 0$. Then

$$\frac{1}{D(\lambda)} = \frac{1}{D_0} \frac{1}{1 + x(\lambda)}.$$

Using

$$\frac{1}{1 + x} = 1 - x + O(x^2) \quad \text{as } x \rightarrow 0,$$

we obtain

$$\frac{1}{1 + x(\lambda)} = 1 - x(\lambda) + O(\lambda^2),$$

and hence

$$\frac{1}{D(\lambda)} = \frac{1}{D_0} - \lambda \frac{D_1}{D_0^2} + O(\lambda^2).$$

With

$$N(\lambda) = N_0 + \lambda N_1 + O(\lambda^2),$$

we have

$$\begin{aligned} \frac{N(\lambda)}{D(\lambda)} &= (N_0 + \lambda N_1 + O(\lambda^2)) \left(\frac{1}{D_0} - \lambda \frac{D_1}{D_0^2} + O(\lambda^2) \right) \\ &= N_0 \cdot \frac{1}{D_0} + \lambda \left(N_1 \cdot \frac{1}{D_0} - N_0 \cdot \frac{D_1}{D_0^2} \right) + O(\lambda^2) \\ &= \frac{N_0}{D_0} + \lambda \frac{N_1 D_0 - N_0 D_1}{D_0^2} + O(\lambda^2), \end{aligned}$$

which is the desired expansion. $\square$

Theorem 1 (First-order expansion of the tilted FM gradient). *Let (X_t, t) denote the space-time pair arising from the rectified flow matching construction, with law $\mathbb{P}_{X_t, t}$ (induced by the joint law of the underlying data, e.g. (X_0, X_1) and the sampling of t), and assume this law is independent of the model parameters θ .*

For each space-time point (x, t) , let z be the conditional path data variable attached to (x, t) , with conditional law $p(z \mid x, t)$, and write

$$U_t(x, z) \in \mathbb{R}^d$$

for the corresponding microscopic target velocity at (x, t) . Let

$$u_\theta^t(x) \in \mathbb{R}^d$$

be a model prediction, differentiable in θ .

Define the conditional covariance and third-order skew moment

$$\Sigma_t(x) := \mathbb{E}_{z \mid x, t} [\varepsilon_t(x, z) \varepsilon_t(x, z)^\top], \quad S_t(x) := \mathbb{E}_{z \mid x, t} [\varepsilon_t(x, z) \|\varepsilon_t(x, z)\|^2].$$

For $\lambda > 0$, define the tilted local loss at (x, t) by

$$\ell_\lambda(x, t; \theta) := \frac{1}{\lambda} \log \mathbb{E}_{z \mid x, t} \exp(\lambda \|u_\theta^t(x) - U_t(x, z)\|^2) \quad (\lambda > 0),$$

with the convention

$$\ell_0(x, t; \theta) := \mathbb{E}_{z \mid x, t} [\|u_\theta^t(x) - U_t(x, z)\|^2].$$

Let the global tilted and FM/MSE objectives be

$$L_\lambda(\theta) := \mathbb{E}_{(X_t, t)} [\ell_\lambda(X_t, t; \theta)], \quad \mathcal{L}_{\text{MSE}}(\theta) := \mathbb{E}_{(X_t, t)} [\ell_0(X_t, t; \theta)].$$

Assume furthermore that

- for each (x, t) , the map $\theta \mapsto u_\theta^t(x)$ is C^1 ,
- there exists an integrable function $G(x, t)$ such that

$$\|\nabla_\theta u_\theta^t(x)\| \leq G(x, t) \quad \text{for all } \theta \text{ in the parameter region of interest,}$$

and

- the conditional moments $\Sigma_t(x)$, $S_t(x)$ and the fourth moments of $\varepsilon_t(x, z)$ are uniformly bounded on the support of (X_t, t) .

Then there exists $\lambda_0 > 0$ and a constant $C < \infty$ such that for all $|\lambda| \leq \lambda_0$,

$$\nabla_\theta L_\lambda(\theta) - \nabla_\theta \mathcal{L}_{\text{MSE}}(\theta) = \tag{31}$$

$$2 \mathbb{E}_{(X_t, t)} \left[\left(-\lambda S_t(X_t) + 2\lambda \Sigma_t(X_t) m_t(X_t) \right)^\top \nabla_\theta u_\theta^t(X_t) \right] + R_\lambda(\theta), \tag{32}$$

where the remainder satisfies

$$\|R_\lambda(\theta)\| \leq C \lambda^2 \quad \text{for all } |\lambda| \leq \lambda_0.$$

Proof of Theorem. Fix (x, t) and write

$$U := U_t(x, z), \quad z \sim p(z \mid x, t),$$

$$u_\theta := u_\theta^t(x),$$

so that U is the random target and u_θ is the model prediction at (x, t) .

For $\lambda \geq 0$ define

$$Z_\lambda(x, t; \theta) := \mathbb{E}_{z|x, t} \exp(\lambda \|u_\theta - U\|^2). \tag{33}$$

The local log-exponential loss is

$$\ell_\lambda(x, t; \theta) := \frac{1}{\lambda} \log Z_\lambda(x, t; \theta) \quad (\lambda > 0), \tag{34}$$

with the convention that ℓ_0 is the usual squared loss. The global loss is

$$L_\lambda(\theta) := \mathbb{E}_{(X_t, t)} [\ell_\lambda(X_t, t; \theta)], \tag{35}$$

where the sampling law of (X_t, t) is fixed and does not depend on θ .

We first compute the gradient of $\ell_\lambda(x, t; \theta)$ at fixed (x, t) . By the chain rule,

$$\nabla_\theta \ell_\lambda(x, t; \theta) = \frac{1}{\lambda} \frac{\nabla_\theta Z_\lambda(x, t; \theta)}{Z_\lambda(x, t; \theta)}. \quad (36)$$

Differentiating (33) under the expectation and using

$$\|u_\theta - U\|^2 = (u_\theta - U)^\top (u_\theta - U), \quad \nabla_\theta \|u_\theta - U\|^2 = 2 (u_\theta - U)^\top \nabla_\theta u_\theta,$$

we obtain

$$\begin{aligned} \nabla_\theta Z_\lambda(x, t; \theta) &= \mathbb{E}_{z|x, t} \left[\exp(\lambda \|u_\theta - U\|^2) \lambda \nabla_\theta \|u_\theta - U\|^2 \right] \\ &= 2\lambda \mathbb{E}_{z|x, t} \left[\exp(\lambda \|u_\theta - U\|^2) (u_\theta - U)^\top \right] \nabla_\theta u_\theta. \end{aligned} \quad (37)$$

Substituting (37) into (36) gives

$$\nabla_\theta \ell_\lambda(x, t; \theta) = 2 \frac{\mathbb{E}_{z|x, t} [\exp(\lambda \|u_\theta - U\|^2) (u_\theta - U)^\top]}{\mathbb{E}_{z|x, t} [\exp(\lambda \|u_\theta - U\|^2)]} \nabla_\theta u_\theta. \quad (38)$$

Introduce the normalized weights

$$\beta_\lambda(z \mid x, t; \theta) := \frac{\exp(\lambda \|u_\theta - U\|^2)}{\mathbb{E}_{z'|x, t} \exp(\lambda \|u_\theta - U_t(x, z')\|^2)}, \quad (39)$$

so that

$$\mathbb{E}_{z|x, t} [\beta_\lambda(z \mid x, t; \theta)] = 1.$$

Then the fraction in (38) can be written as

$$\mathbb{E}_{z|x, t} [\beta_\lambda(z \mid x, t; \theta) (u_\theta - U)] = u_\theta - \mathbb{E}_{z|x, t} [\beta_\lambda(z \mid x, t; \theta) U].$$

We define the tilted conditional mean

$$m_\lambda^t(x; \theta) := \mathbb{E}_{z|x, t} [\beta_\lambda(z \mid x, t; \theta) U_t(x, z)], \quad (40)$$

so that

$$\mathbb{E}_{z|x, t} [\beta_\lambda(z \mid x, t; \theta) (u_\theta - U)] = u_\theta - m_\lambda^t(x; \theta).$$

Equation (38) then becomes

$$\nabla_\theta \ell_\lambda(x, t; \theta) = 2 (u_\theta - m_\lambda^t(x; \theta))^\top \nabla_\theta u_\theta. \quad (41)$$

Taking the expectation in (35) and using that the distribution of (X_t, t) does not depend on θ , we obtain the global gradient identity

$$\nabla_{\theta} L_{\lambda}(\theta) = 2 \mathbb{E}_{(X_t, t)} \left[(u_{\theta}^t(X_t) - m_{\lambda}^t(X_t; \theta))^{\top} \nabla_{\theta} u_{\theta}^t(X_t) \right]. \quad (42)$$

For comparison, the FM/MSE objective has gradient

$$\nabla_{\theta} \mathcal{L}_{\text{MSE}}(\theta) = 2 \mathbb{E}_{(X_t, t)} \left[(u_{\theta}^t(X_t) - \mu_t(X_t))^{\top} \nabla_{\theta} u_{\theta}^t(X_t) \right], \quad (43)$$

where $\mu_t(x) := \mathbb{E}_{z|x,t}[U_t(x, z)]$ is the usual conditional mean. Subtracting (43) from (42) and rearranging gives the exact difference

$$\nabla_{\theta} L_{\lambda}(\theta) - \nabla_{\theta} \mathcal{L}_{\text{MSE}}(\theta) = 2 \mathbb{E}_{(X_t, t)} \left[(\mu_t(X_t) - m_{\lambda}^t(X_t; \theta))^{\top} \nabla_{\theta} u_{\theta}^t(X_t) \right]. \quad (44)$$

Moment expansion of the tilted mean. We now expand $m_{\lambda}^t(x; \theta)$ in powers of λ at fixed (x, t) . Write

$$U_t(x, z) = \mu_t(x) + \varepsilon, \quad \mathbb{E}[\varepsilon \mid x, t] = 0,$$

and define the conditional moments

$$\Sigma_t(x) = \mathbb{E}[\varepsilon \varepsilon^{\top} \mid x, t], \quad S_t(x) = \mathbb{E}[\varepsilon \|\varepsilon\|^2 \mid x, t].$$

Let the residual error be $m_t(x) = u_{\theta}^t(x) - \mu_t(x)$; then

$$\|u_{\theta}^t(x) - U_t(x, z)\|^2 = \|m_t(x) - \varepsilon\|^2 = \|m_t(x)\|^2 + \|\varepsilon\|^2 - 2 m_t(x)^{\top} \varepsilon. \quad (45)$$

Introduce

$$a := \|\varepsilon\|^2 - 2 m_t(x)^{\top} \varepsilon. \quad (46)$$

The common factor $\exp(\lambda \|m_t(x)\|^2)$ cancels when forming the normalized weights, so from (40) we may equivalently write

$$m_{\lambda}^t(x; \theta) = \frac{\mathbb{E}_{z|x,t}[(\mu_t(x) + \varepsilon) e^{\lambda a}]}{\mathbb{E}_{z|x,t}[e^{\lambda a}]}. \quad (47)$$

For small λ we expand

$$e^{\lambda a} = 1 + \lambda a + O(\lambda^2), \quad (48)$$

where $O(\lambda^2)$ denotes terms whose norm is bounded by a constant times λ^2 for λ in a fixed small interval; in particular such terms vanish faster than λ as $\lambda \downarrow 0$.

Using (48), the numerator

$$N(\lambda) := \mathbb{E}_{z|x,t}[(\mu_t(x) + \varepsilon) e^{\lambda a}]$$

satisfies

$$\begin{aligned} N(\lambda) &= \mathbb{E}[\mu_t(x) + \varepsilon] + \lambda \mathbb{E}[(\mu_t(x) + \varepsilon)a] + O(\lambda^2) \\ &= \mu_t(x) + \lambda(\mu_t(x) \mathbb{E}[a] + \mathbb{E}[\varepsilon a]) + O(\lambda^2), \end{aligned}$$

since $\mathbb{E}[\varepsilon | x, t] = 0$. Similarly, the denominator

$$D(\lambda) := \mathbb{E}_{z|x,t}[e^{\lambda a}]$$

has the expansion

$$D(\lambda) = 1 + \lambda \mathbb{E}[a] + O(\lambda^2).$$

We can now use the elementary first-order ratio expansion

$$\frac{N_0 + \lambda N_1 + O(\lambda^2)}{D_0 + \lambda D_1 + O(\lambda^2)} = \frac{N_0}{D_0} + \lambda \frac{N_1 D_0 - N_0 D_1}{D_0^2} + O(\lambda^2), \quad (49)$$

which follows by expanding $1/(D_0 + \lambda D_1)$ to first-order and multiplying out. In our case,

$$N_0 = \mu_t(x), \quad N_1 = \mu_t(x) \mathbb{E}[a] + \mathbb{E}[\varepsilon a], \quad D_0 = 1, \quad D_1 = \mathbb{E}[a],$$

so (49) yields

$$\begin{aligned} m_\lambda^t(x; \theta) &= \mu_t(x) + \lambda(N_1 - \mu_t(x) D_1) + O(\lambda^2) \\ &= \mu_t(x) + \lambda \mathbb{E}[\varepsilon a] + O(\lambda^2). \end{aligned} \quad (50)$$

From (46),

$$\varepsilon a = \varepsilon \|\varepsilon\|^2 - 2 \varepsilon \varepsilon^\top m_t(x),$$

so

$$\begin{aligned} \mathbb{E}_{z|x,t}[\varepsilon a] &= \mathbb{E}[\varepsilon \|\varepsilon\|^2 | x, t] - 2 \mathbb{E}[\varepsilon \varepsilon^\top | x, t] m_t(x) \\ &= S_t(x) - 2 \Sigma_t(x) m_t(x). \end{aligned}$$

Substituting into (50) gives the moment expansion

$$m_\lambda^t(x; \theta) = \mu_t(x) + \lambda S_t(x) - 2\lambda \Sigma_t(x) m_t(x) + O(\lambda^2). \quad (51)$$

First-order correction to the gradient. Finally we combine (51) with the exact gradient difference (44). At fixed (x, t) ,

$$\mu_t(x) - m_\lambda^t(x; \theta) = -\lambda S_t(x) + 2\lambda \Sigma_t(x) m_t(x) + O(\lambda^2),$$

and so

$$\begin{aligned} \nabla_\theta L_\lambda(\theta) - \nabla_\theta \mathcal{L}_{\text{MSE}}(\theta) &= 2 \mathbb{E}_{(X_t, t)} \left[(\mu_t(X_t) - m_\lambda^t(X_t; \theta))^\top \nabla_\theta u_\theta^t(X_t) \right] \\ &= 2 \mathbb{E}_{(X_t, t)} \left[(-\lambda S_t(X_t) + 2\lambda \Sigma_t(X_t) m_t(X_t))^\top \nabla_\theta u_\theta^t(X_t) \right] + O(\lambda^2). \end{aligned} \quad (52)$$

Thus, for small λ , the entropic FM gradient is the FM/MSE gradient plus a first-order correction consisting of:

- a covariance-preconditioning term $2\lambda \Sigma_t(x) m_t(x)$, and
- a skew/tail term $-\lambda S_t(x)$ that depends only on the data moments and does not involve the current residual.

The $O(\lambda^2)$ remainder term collects all higher-order contributions and vanishes faster than λ as λ approaches 0. □

Gradient transfer to the marginal FM objective. Define the marginal FM loss

$$L_{\text{FM}}(\theta) = \mathbb{E} [\|u_\theta^t(X_t) - u_t(X_t)\|^2],$$

with conditional mean

$$\mu_t(x) = \mathbb{E}_{z|x, t} [U_t(x, z)].$$

Corollary 1. Immediate corollary (by substitution). As shown in [1],

$$\nabla_\theta L_{\text{FM}}(\theta) = \nabla_\theta L_{\text{MSE}}(\theta),$$

which together with 52 implies the same first-order correction applies to the marginal FM gradient:

$$\nabla_\theta L_{\text{FM}}(\theta) \longrightarrow \nabla_\theta L_{\text{FM}}(\theta) + 2 \mathbb{E} \left[(-\lambda S_t(X_t) + 2\lambda \Sigma_t(X_t) m_t(X_t))^\top \nabla_\theta u_\theta^t(X_t) \right].$$

□

Remarkably, definitions of S and m_t imply that if the local velocity fields are symmetric and the baseline model perfectly estimates the FM velocity field, the first order correction terms disappear; i.e, The *original* FM optimum is a fixed point of our risk-sensitive / entropic deformation under those assumptions.