

An introduction to generative flow models

From flow matching to diffusion

Julian Brandon - Les Gustins 2026



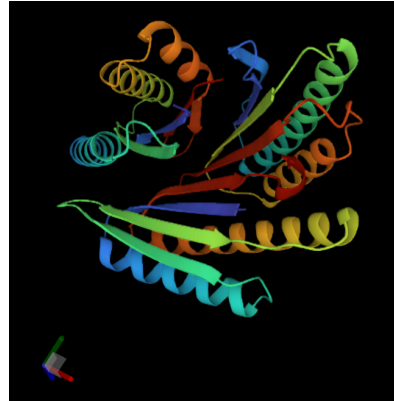
Neural networks and generative AI

◆ Image generation



Stable diffusion XL

◆ Protein modelling



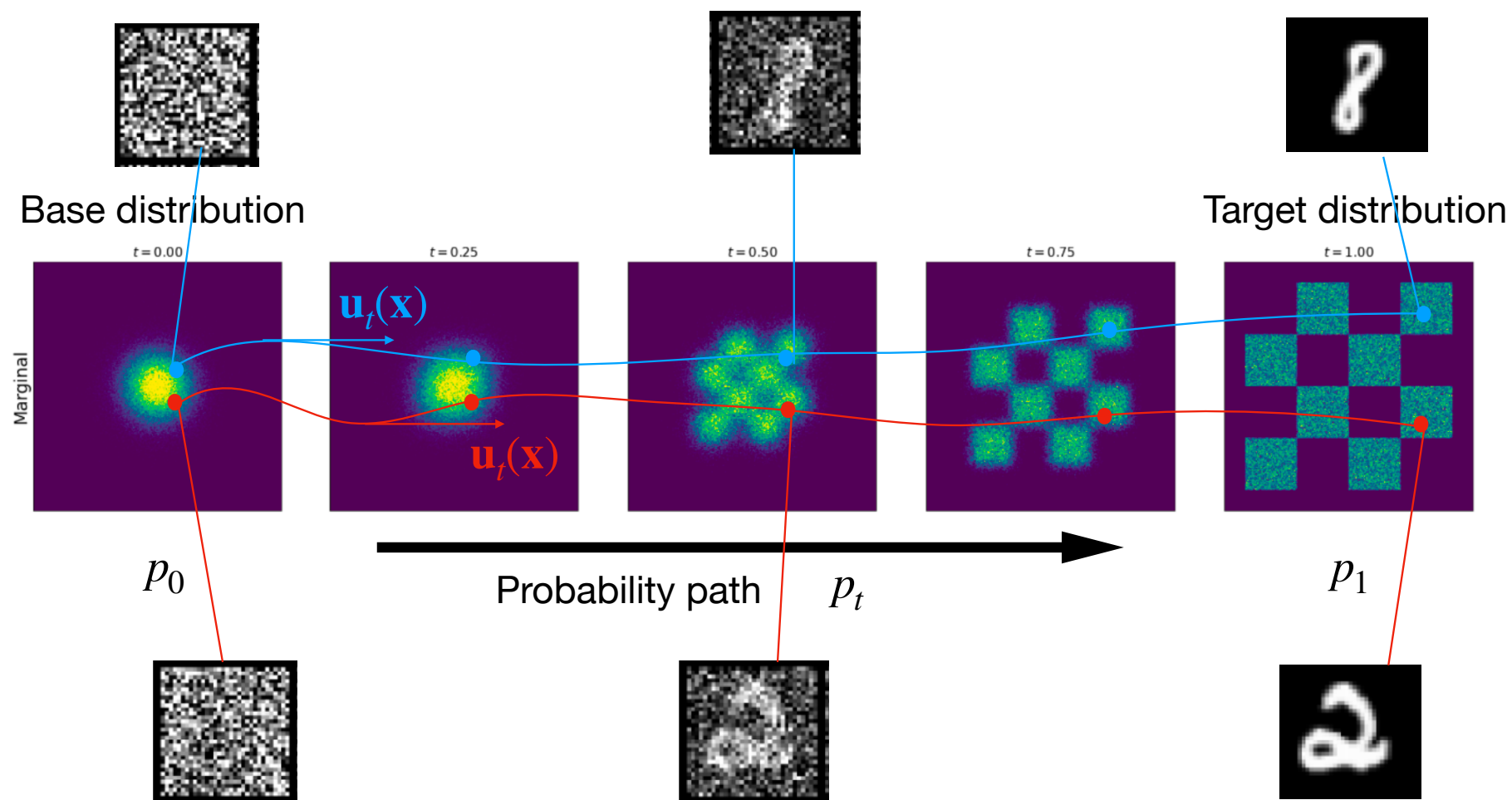
Chroma

◆ Robotics



Diffusion CCSP

Generative flow models

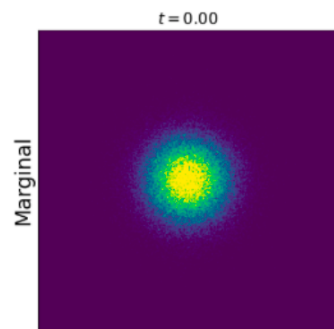


Questions

1. What should be the form of p_t and $\mathbf{u}_t$?
2. How to train a neural network, $\mathbf{u}_t^\theta$, to approximate $\mathbf{u}_t$?
3. How can we extend to stochastic dynamics?

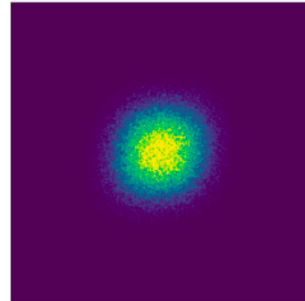
Constructing a probability path

Base distribution

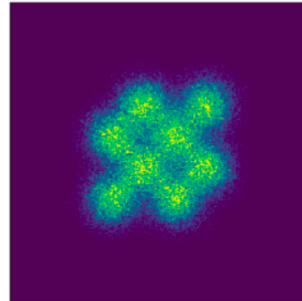


p_0

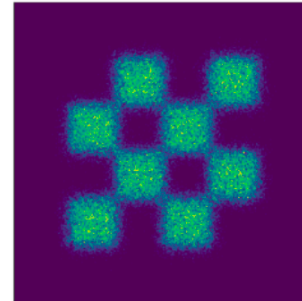
$t = 0.25$



$t = 0.50$

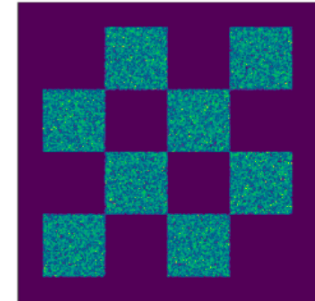


$t = 0.75$



Target distribution

$t = 1.00$



p_1

Probability path

p_t

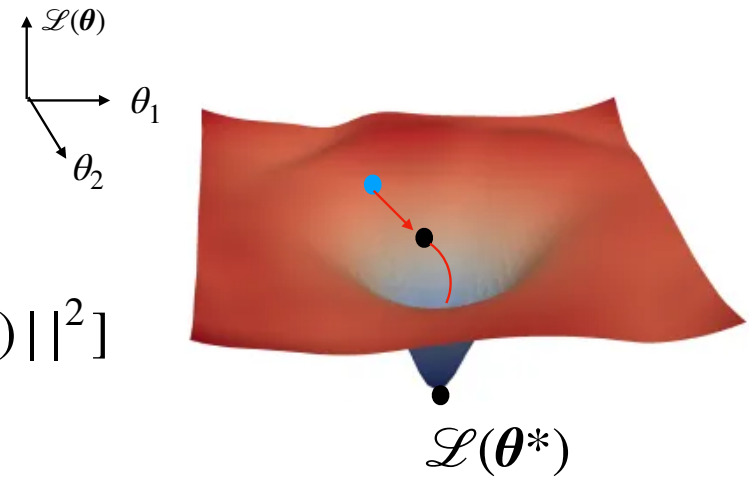
Training a neural network with gradient descent

Neural network: $\mathbf{u}_t^\theta(\mathbf{x})$

Target velocity field: $\mathbf{u}_t^{\text{target}}(\mathbf{x})$

Objective fn: $\mathcal{L}(\theta) = \mathbb{E}_{\mathbf{x} \sim p_{t,t}}[||\mathbf{u}_t^\theta(\mathbf{x}) - \mathbf{u}_t^{\text{target}}(\mathbf{x})||^2]$

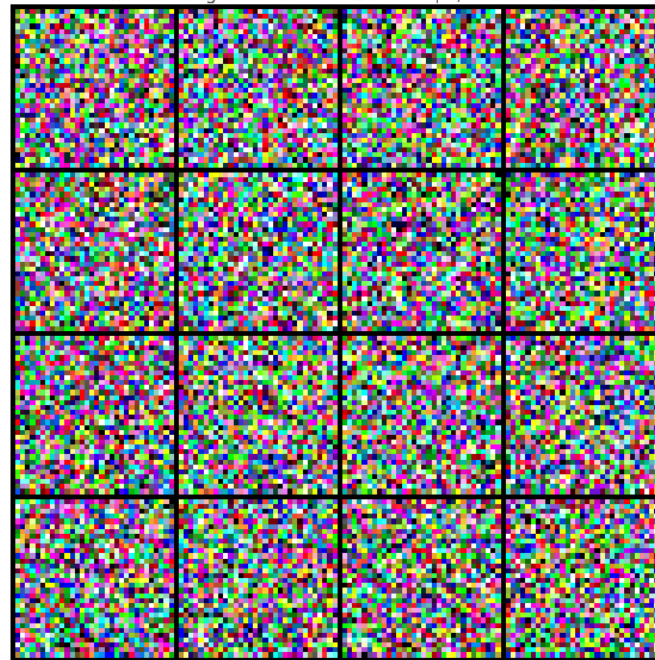
Parameter update: $\theta_{\tau+1} = \theta_\tau - \eta \nabla_\theta \mathcal{L}(\theta_\tau)$



Flow matching

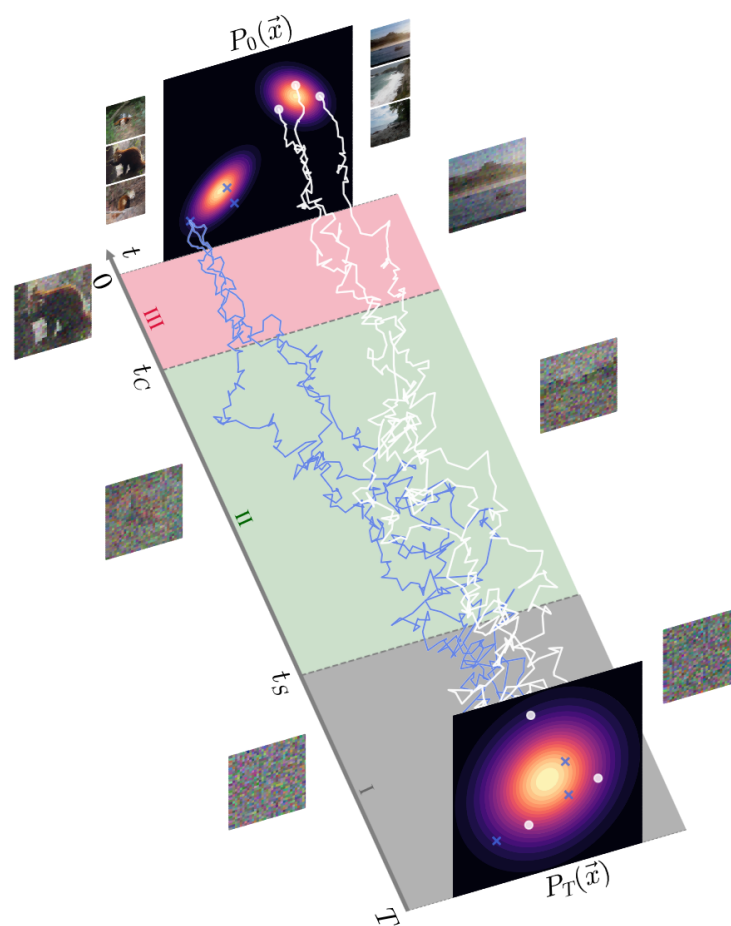
t = 0.00 step 0/60

Flow matching inference: 60 Euler steps, CelebA 32x32



$x \leftarrow x + v(x, t) * dt$

Diffusion



Biroli et al. 2025

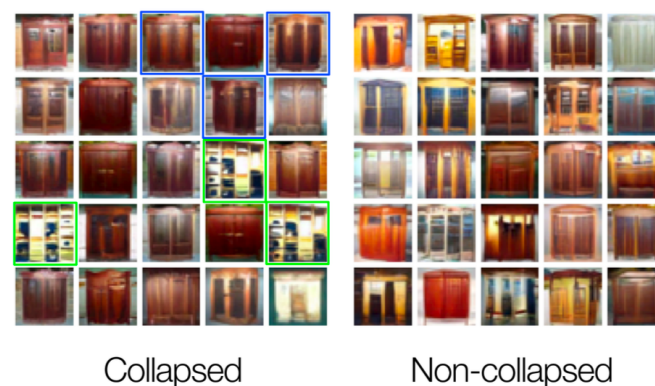
Generative models show biases

Memorisation of data samples



Gu et al. 2025 TMLR

Mode collapse



Qin et al. 2023, CVPR

These models are in principle *expressive enough* to learn the right data features. This is a bias coming from the learning dynamics.

An operator view of gradient flow

$$\mathcal{L}(W) = \mathbb{E}_{\mathbf{x}_0, \mathbb{B}} \left[\int_0^1 l(\mathbf{x}(t), t) dt \right]$$

Theorem (Operator View of Gradient)

The gradient of $\mathcal{L}$ with respect to a linear parameter $W \in \mathbb{R}^{m \times k}$ can be written as the composition of two linear operators:

$$\nabla_W \mathcal{L} = A \circ X$$

$$A : L^2(\mathbb{R}) \rightarrow \mathbb{R}^m, \quad X : \mathbb{R}^k \rightarrow L^2(\mathbb{R})$$

Theorem (Stable rank bound)

The stable rank of the gradient can be bounded as:

$$\text{SRank}(\nabla_W \mathcal{L}) \leq \frac{\|A\|_2 \|X\|_2}{\|A \circ X\|_2} \min \{ \text{SRank}(A), \text{SRank}(X) \}$$

where A and X are the adjoint and state operators.

Next, we consider the flow-matching objective:

$$\mathcal{L}_{\text{FM}} = \mathbb{E}_{\mathbf{x}_0} \left[\int_0^1 ||\mathbf{u}^*(\mathbf{x}^*(t), t) - \mathbf{u}_\theta(\mathbf{x}^*(t), t)||^2 dt \right]$$

For our linear model that is:

$$\mathcal{L}_{\text{FM}} = \mathbb{E}_{\mathbf{x}_0} \left[\int_0^1 ||\mathcal{W}^* \mathbf{x}^*(t) - \mathcal{W} \mathbf{x}^*(t)||^2 dt \right]$$

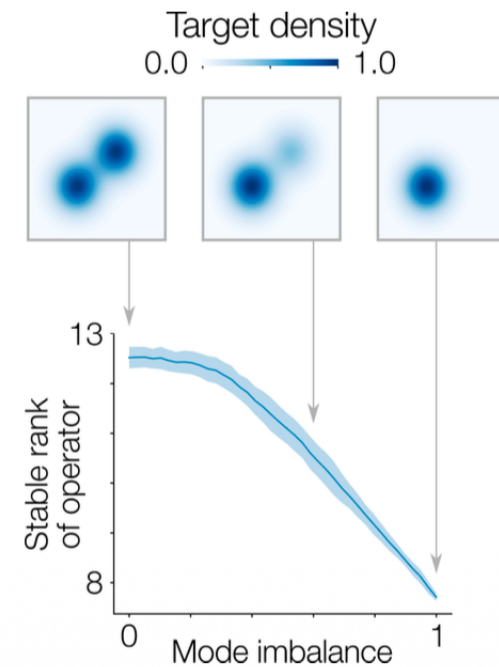
- *Data*: We will consider unequal weights in our Gaussian mixture:

$$p^*(\mathbf{x}) = \sum_{i=1}^m \pi_i \mathcal{N}(\mathbf{x}; \boldsymbol{\mu}_i, \sigma_i I)$$

- *Network*: We will add an additional layer to our neural network:

$$\mathbf{u}(\mathbf{x}^*(t)) = W_2 W_1 \mathbf{x}^*(t)$$

In that case, the operator rank continuously varies with the mixture weights.



Theorem (Informal)

The i 'th mode takes $O(\pi_i^{-2})$ iterations to be learned.

