

Deep Generative Models

12. Score-Based Models



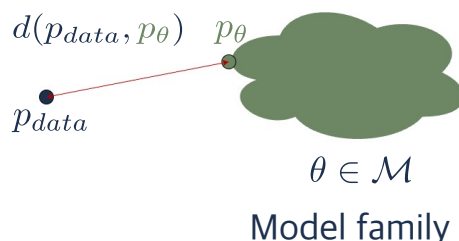
- 국가수리과학연구소 산업수학혁신센터 김민중

Summary

- **Representation:** how do we model the joint distribution of many random variables?
 - Need compact representation
- **Learning:** what is the right way to compare probability distributions?



$$\mathbf{x}_i \sim p_{data}$$
$$i = 1, 2, \dots, N$$



- **Inference:** how do we invert the generation process

Representation

- Probability density function (p.d.f.) or probability mass function (p.m.f.)

$$p(\mathbf{x})$$

- Autoregressive Models

$$p_{\theta}(\mathbf{x}) = \prod_{i=1}^d p_{\theta}(x_i | \mathbf{x}_{<i})$$

- Variational Autoencoders

$$p_{\theta}(\mathbf{x}) = \int p_{\theta}(\mathbf{x}, \mathbf{z}) d\mathbf{z} = \int p_{\theta}(\mathbf{x} | \mathbf{z}) p(\mathbf{z}) d\mathbf{z}$$

Representation

- Probability density function (p.d.f.) or probability mass function (p.m.f.)

$$p(\mathbf{x})$$

- Flow Models

$$p_X(\mathbf{x}; \theta) = p_Z\left(\mathbf{f}_\theta^{-1}(\mathbf{x})\right) \left| \det \left(\frac{\partial \mathbf{f}_\theta^{-1}(\mathbf{x})}{\partial \mathbf{x}} \right) \right|$$

- Energy-based models

$$p_\theta(\mathbf{x}) = \frac{1}{Z(\theta)} \exp(f_\theta(\mathbf{x}))$$

Representation

- Probability density function (p.d.f.) or probability mass function (p.m.f.)

$$p(\boldsymbol{x})$$

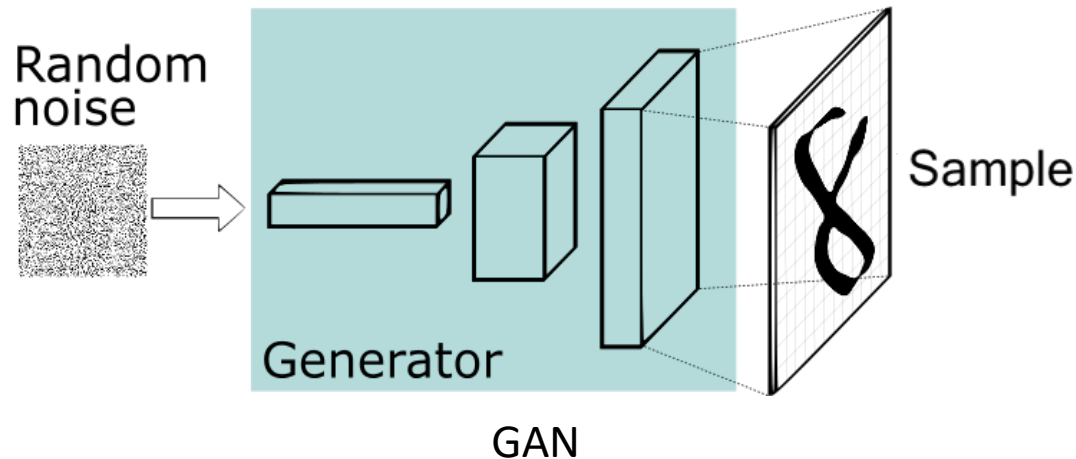
- Pros
 - Maximum likelihood training
 - Principled model comparison via likelihoods
- Cons
 - Special architectures or surrogate losses to deal with intractable partition functions

Implicit generative model

- Generative adversarial networks (GANs)

$$z \sim p(z)$$
$$x = G_{\theta}(z)$$

- A two-player minimax game between a generator and a discriminator(two sample test)



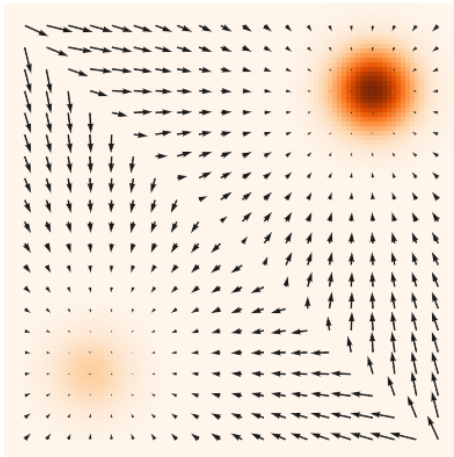
Generative adversarial networks

- Pros
 - Samples typically have better quality
- Cons
 - Require adversarial training
 - Training instability and mode collapse
 - No principled way to compare different models
 - No principled termination criteria for training

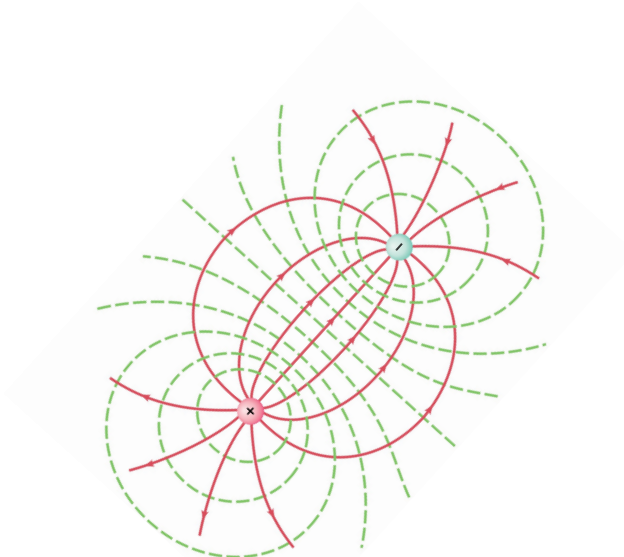
Score functions

- When the pdf is differentiable, we can compute the gradient of a probability density
- Score function

$$\nabla_x \log p(\mathbf{x})$$



(pdf and score)

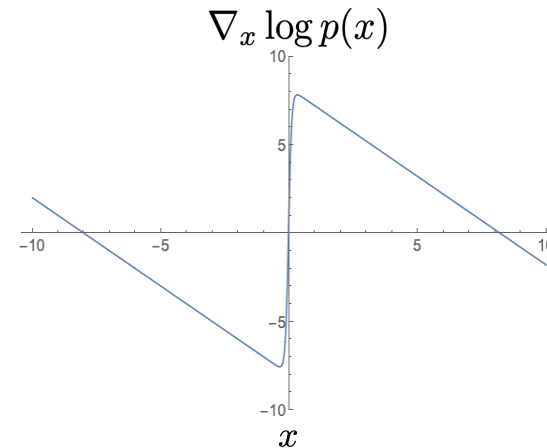
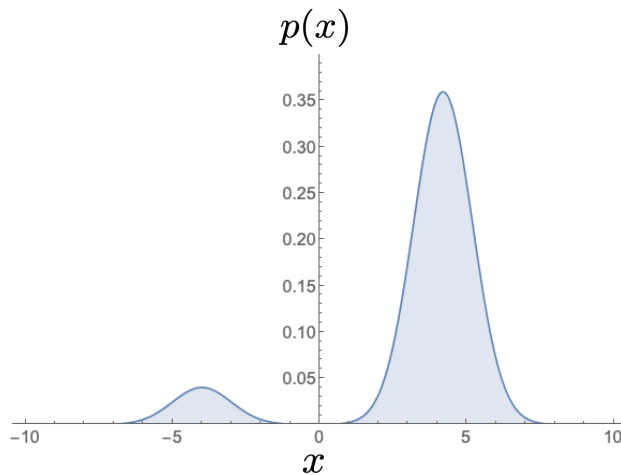


(Electrical potentials and fields)

Score functions

- When the pdf is differentiable, we can compute the gradient of a probability density
- Score function

$$\nabla_x \log p(x)$$



Recap. of Energy-based model

- Energy-based models: $\frac{1}{Z(\theta)} \exp(f_{\theta}(\mathbf{x}))$
 - $Z(\theta)$ is intractable, so no access to likelihood
 - Comparing the probability of two points is easy

$$\frac{p_{\theta}(\mathbf{x})}{p_{\theta}(\mathbf{x}')} = \exp(f_{\theta}(\mathbf{x}) - f_{\theta}(\mathbf{x}'))$$

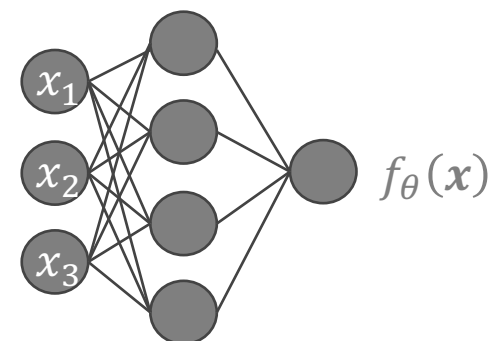
- Maximum likelihood training:

$$\max_{\theta} [f_{\theta}(\mathbf{x}_{train}) - \log Z(\theta)]$$

- Contrastive divergence:

$$\nabla_{\theta} f_{\theta}(\mathbf{x}_{train}) - \nabla_{\theta} \log Z(\theta) \approx \nabla_{\theta} f_{\theta}(\mathbf{x}_{train}) - \nabla_{\theta} f_{\theta}(\mathbf{x}_{sample})$$

- where $\mathbf{x}_{sample} \sim p_{\theta}(\mathbf{x}) = \frac{\exp(f_{\theta}(\mathbf{x}))}{Z(\theta)}$
- Requires iterative sampling during training

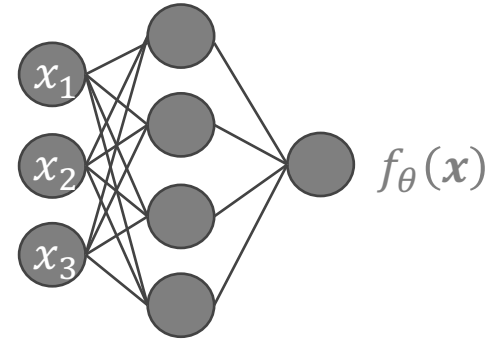


Recap. of Energy-based model

- Energy-based models

$$p_{\theta}(\mathbf{x}) = \frac{1}{Z(\theta)} \exp(f_{\theta}(\mathbf{x}))$$

$f_{\theta}(\mathbf{x}) \in \mathbb{R}$



- Minimizing Fisher divergence

$$\min_{\theta} \frac{1}{2} E_{\mathbf{x} \sim p_{data}} [\|\nabla_{\mathbf{x}} \log p_{data}(\mathbf{x}) - \nabla_{\mathbf{x}} \log p_{\theta}(\mathbf{x})\|_2^2]$$

- Score matching

$$\begin{aligned} & \frac{1}{2} E_{\mathbf{x} \sim p_{data}} [\|\nabla_{\mathbf{x}} \log p_{data}(\mathbf{x}) - \nabla_{\mathbf{x}} \log p_{\theta}(\mathbf{x})\|_2^2] \\ &= E_{\mathbf{x} \sim p_{data}} \left[\frac{1}{2} \|\nabla_{\mathbf{x}} \log p_{\theta}(\mathbf{x})\|_2^2 + \text{tr}(\nabla_{\mathbf{x}}^2 \log p_{\theta}(\mathbf{x})) \right] + \text{const.} \end{aligned}$$

Score matching for training EBMs

- Score function of EBMs

$$\nabla_{\mathbf{x}} \log p_{\theta}(\mathbf{x}) = \nabla_{\mathbf{x}} f_{\theta}(\mathbf{x}) - \nabla_{\mathbf{x}} \log Z(\theta) = \nabla_{\mathbf{x}} f_{\theta}(\mathbf{x})$$

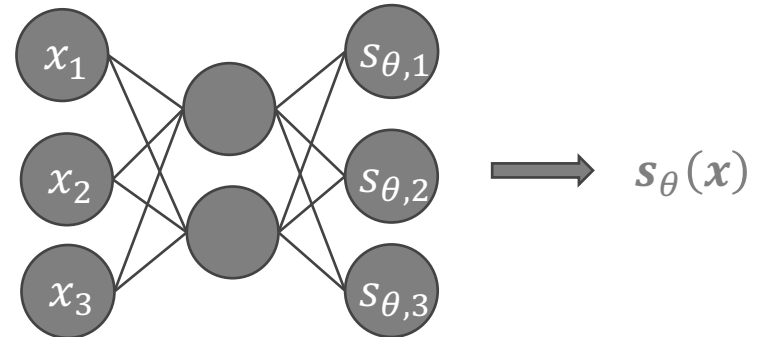
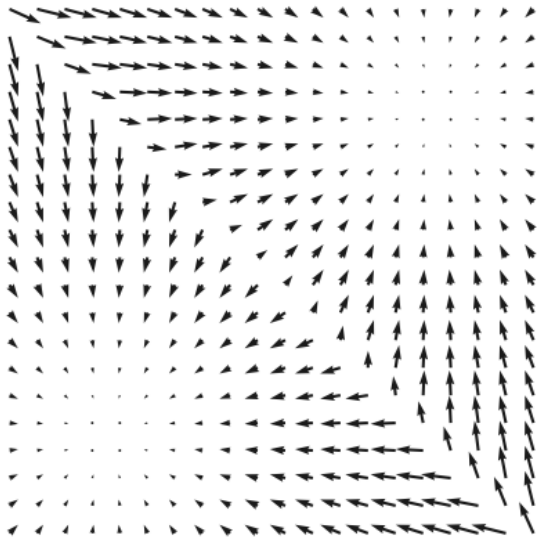
- Score matching for EBMs

$$\begin{aligned} E_{\mathbf{x} \sim p_{data}} \left[\frac{1}{2} \|\nabla_{\mathbf{x}} \log p_{\theta}(\mathbf{x})\|_2^2 + \text{tr}(\nabla_{\mathbf{x}}^2 \log p_{\theta}(\mathbf{x})) \right] \\ = E_{\mathbf{x} \sim p_{data}} \left[\frac{1}{2} \|\nabla_{\mathbf{x}} f_{\theta}(\mathbf{x})\|_2^2 + \text{tr}(\nabla_{\mathbf{x}}^2 f_{\theta}(\mathbf{x})) \right] \end{aligned}$$

Score-based models

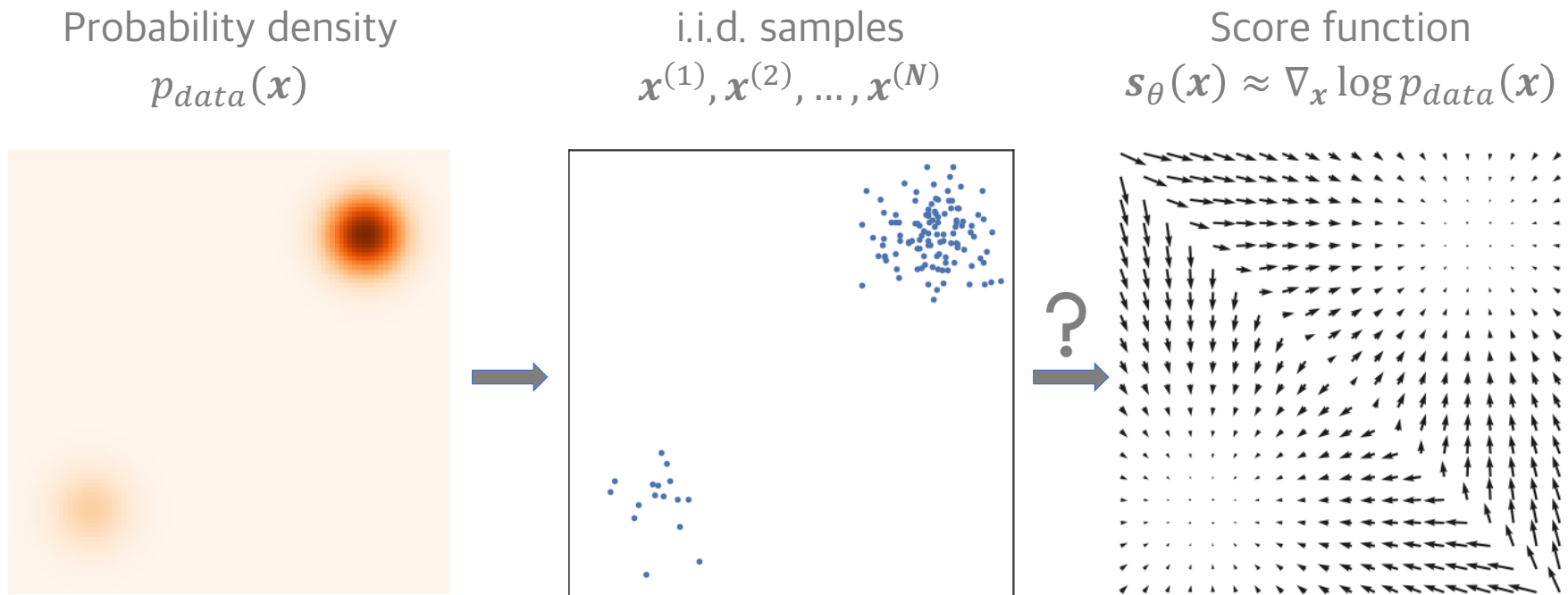
- Directly model the vector field of gradients

$$\begin{aligned} s_{\theta}(\mathbf{x}) &: \mathbb{R}^d \rightarrow \mathbb{R}^d \\ s_{\theta}(\mathbf{x}) &\approx \nabla_{\mathbf{x}} \log p_{data}(\mathbf{x}) \end{aligned}$$



Framework for score estimation

- Score estimation by training score-based models



Score estimation by training score-based models

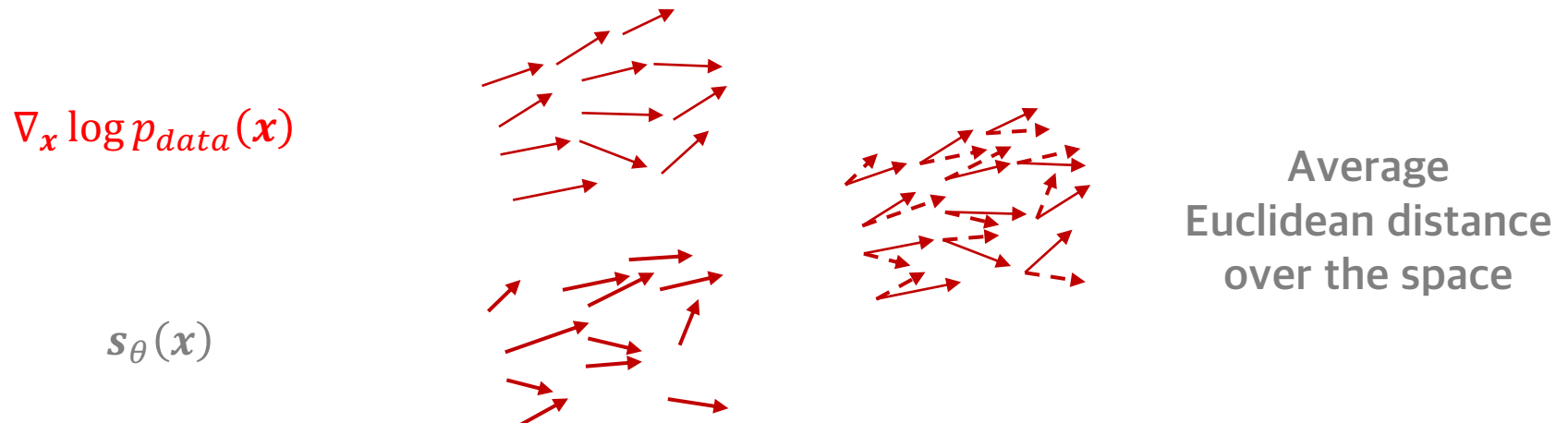
- **Given:** i.i.d. samples $\{\mathbf{x}^{(1)}, \mathbf{x}^{(2)}, \dots, \mathbf{x}^{(N)}\} \sim p_{data}(\mathbf{x})$
- **Task:** Estimating the score $\nabla_{\mathbf{x}} \log p_{data}(\mathbf{x})$
- **Score model:** A learnable vector valued function

$$\mathbf{s}_{\theta}(\mathbf{x}): \mathbb{R}^d \rightarrow \mathbb{R}^d$$

- **Goal**

$$\mathbf{s}_{\theta}(\mathbf{x}) \approx \nabla_{\mathbf{x}} \log p_{data}(\mathbf{x})$$

- How to compare two vector fields of scores?



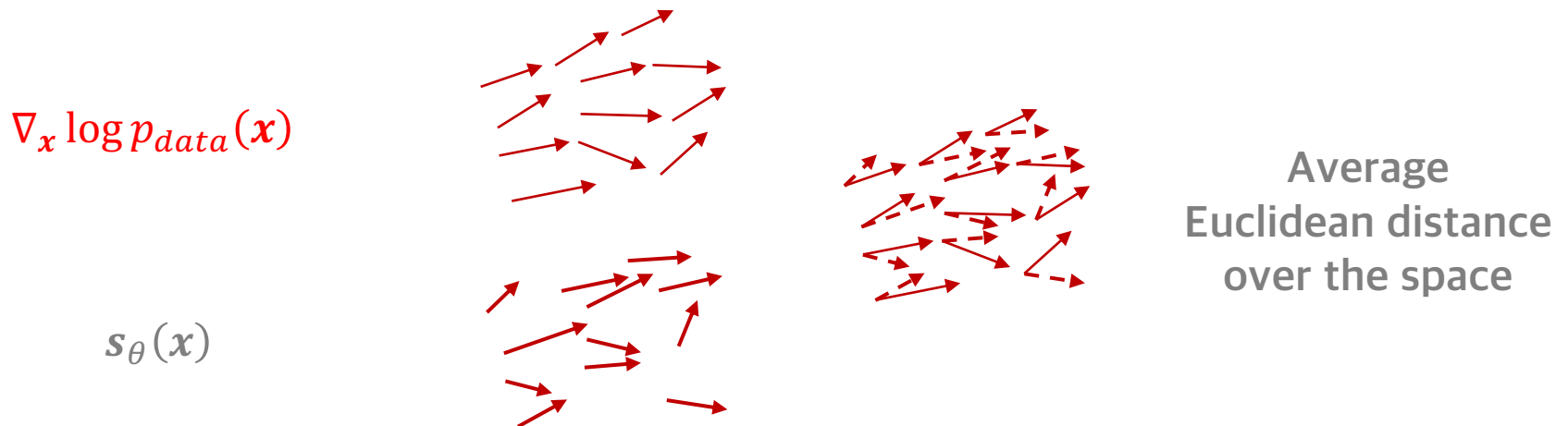
Score estimation by training score-based models

- **Objective:** Average Euclidean distance over the whole space

$$\frac{1}{2} E_{\mathbf{x} \sim p_{data}} [\|\nabla_{\mathbf{x}} \log p_{data}(\mathbf{x}) - \mathbf{s}_{\theta}(\mathbf{x})\|_2^2]$$

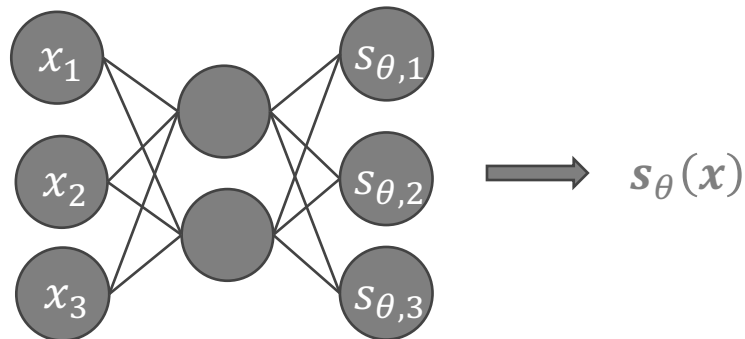
- **Score matching**

$$E_{\mathbf{x} \sim p_{data}} \left[\frac{1}{2} \|\mathbf{s}_{\theta}(\mathbf{x})\|_2^2 + \text{tr}(\nabla_{\mathbf{x}} \mathbf{s}_{\theta}(\mathbf{x})) \right]$$



Score matching is not scalable

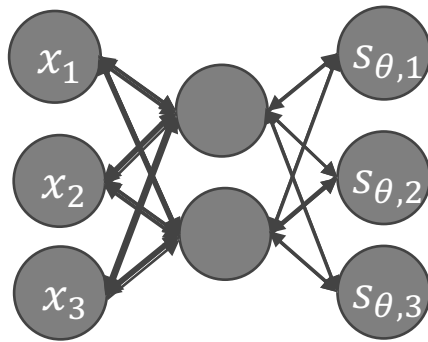
- Deep neural networks as more expressive score models



Not scalable!

- Need to compute $\|s_{\theta}(x)\|_2^2$ and $\text{tr}(\nabla_x s_{\theta}(x))$

$$\begin{pmatrix} \frac{\partial s_{\theta,1}(x)}{\partial x_1} \\ \frac{\partial s_{\theta,2}(x)}{\partial x_2} \\ \frac{\partial s_{\theta,3}(x)}{\partial x_3} \end{pmatrix}$$



$$s_{\theta,1}(x)$$

$$s_{\theta,2}(x)$$

$$s_{\theta,3}(x)$$

$$\nabla_x s_{\theta}(x) = \begin{pmatrix} \frac{\partial s_{\theta,1}(x)}{\partial x_1} & \frac{\partial s_{\theta,1}(x)}{\partial x_2} & \frac{\partial s_{\theta,1}(x)}{\partial x_3} \\ \frac{\partial s_{\theta,2}(x)}{\partial x_1} & \frac{\partial s_{\theta,2}(x)}{\partial x_2} & \frac{\partial s_{\theta,2}(x)}{\partial x_3} \\ \frac{\partial s_{\theta,3}(x)}{\partial x_1} & \frac{\partial s_{\theta,3}(x)}{\partial x_2} & \frac{\partial s_{\theta,3}(x)}{\partial x_3} \end{pmatrix}$$

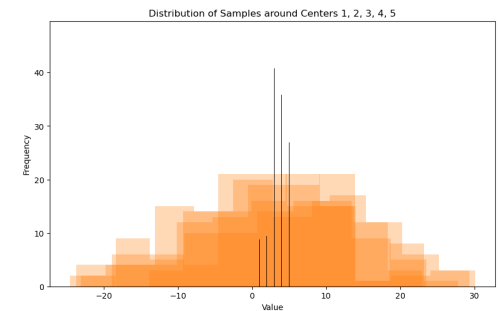
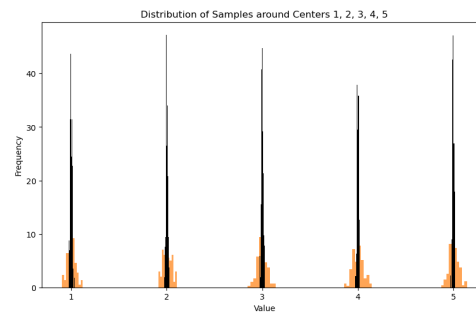
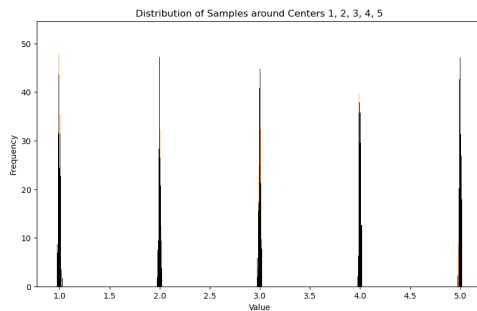
Scalable approaches to score estimation

- Denoising score matching (Vincent, 2011)
- Sliced score matching (Song & Garg, 2019)
- Need to eliminate the trace of Jacobian term in the training objective

Denoising score matching (Vincent, 2011)

- Consider the perturbed distribution

$$q_{\sigma}(\tilde{\mathbf{x}}|\mathbf{x}) := N(\tilde{\mathbf{x}}|\mathbf{x}, \sigma^2 \mathbf{I}), \quad q_{\sigma}(\tilde{\mathbf{x}}) = \int p_{data}(\mathbf{x}) q_{\sigma}(\tilde{\mathbf{x}}|\mathbf{x}) d\mathbf{x}$$



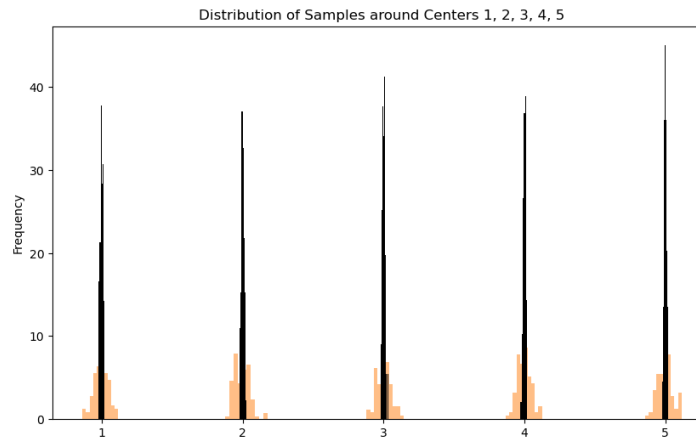
$p_{data}(\mathbf{x})$

$q_{\sigma}(\mathbf{x})$

Denoising score matching (Vincent, 2011)

- Consider the perturbed distribution

$$q_{\sigma}(\tilde{\mathbf{x}}|\mathbf{x}) := N(\tilde{\mathbf{x}}|\mathbf{x}, \sigma^2 \mathbf{I}), \quad q_{\sigma}(\tilde{\mathbf{x}}) = \int p_{data}(\mathbf{x}) q_{\sigma}(\tilde{\mathbf{x}}|\mathbf{x}) d\mathbf{x}$$



$p_{data}(\mathbf{x})$

$q_{\sigma}(\mathbf{x})$

- Score estimation for $\nabla_{\tilde{\mathbf{x}}} \log q_{\sigma}(\tilde{\mathbf{x}})$ is easier
- If the noise level is small, this is a good approximation $q_{\sigma} \approx p_{data}$

Denoising score matching

- Matching the score of a noise-perturbed distribution
- Denoising score matching: $s_{\theta}(\tilde{\mathbf{x}}) \approx \nabla_{\tilde{\mathbf{x}}} \log q_{\sigma}(\tilde{\mathbf{x}})$
$$\frac{1}{2} E_{\tilde{\mathbf{x}} \sim q_{\sigma}} [\|\nabla_{\tilde{\mathbf{x}}} \log q_{\sigma}(\tilde{\mathbf{x}}) - \mathbf{s}_{\theta}(\tilde{\mathbf{x}})\|_2^2]$$
- However, $q_{\sigma}(\tilde{\mathbf{x}}) = \int p_{data}(\mathbf{x}) q_{\sigma}(\tilde{\mathbf{x}}|\mathbf{x}) d\mathbf{x}$ is unknown

Denoising score matching

- Fisher divergence

$$\begin{aligned} & \frac{1}{2} E_{\tilde{\mathbf{x}} \sim q_{\sigma}} [\|\nabla_{\tilde{\mathbf{x}}} \log q_{\sigma}(\tilde{\mathbf{x}}) - \mathbf{s}_{\theta}(\tilde{\mathbf{x}})\|_2^2] \\ &= \frac{1}{2} \int q_{\sigma}(\tilde{\mathbf{x}}) \|\nabla_{\tilde{\mathbf{x}}} \log q_{\sigma}(\tilde{\mathbf{x}}) - \mathbf{s}_{\theta}(\tilde{\mathbf{x}})\|_2^2 d\tilde{\mathbf{x}} \\ &= \frac{1}{2} \int q_{\sigma}(\tilde{\mathbf{x}}) \|\nabla_{\tilde{\mathbf{x}}} \log q_{\sigma}(\tilde{\mathbf{x}})\|_2^2 d\tilde{\mathbf{x}} + \frac{1}{2} \int q_{\sigma}(\tilde{\mathbf{x}}) \|\mathbf{s}_{\theta}(\tilde{\mathbf{x}})\|_2^2 d\tilde{\mathbf{x}} \\ &\quad - \int q_{\sigma}(\tilde{\mathbf{x}}) \nabla_{\tilde{\mathbf{x}}} \log q_{\sigma}(\tilde{\mathbf{x}})^T \mathbf{s}_{\theta}(\tilde{\mathbf{x}}) d\tilde{\mathbf{x}} \\ &= \text{const.} + \frac{1}{2} \int q_{\sigma}(\tilde{\mathbf{x}}) \|\mathbf{s}_{\theta}(\tilde{\mathbf{x}})\|_2^2 d\tilde{\mathbf{x}} - \int q_{\sigma}(\tilde{\mathbf{x}}) \nabla_{\tilde{\mathbf{x}}} \log q_{\sigma}(\tilde{\mathbf{x}})^T \mathbf{s}_{\theta}(\tilde{\mathbf{x}}) d\tilde{\mathbf{x}} \end{aligned}$$

Denoising score matching

$$\begin{aligned}\frac{1}{2} \int q_{\sigma}(\tilde{\mathbf{x}}) \|\mathbf{s}_{\theta}(\tilde{\mathbf{x}})\|_2^2 d\tilde{\mathbf{x}} &= \frac{1}{2} \int \int p_{data}(\mathbf{x}) q_{\sigma}(\tilde{\mathbf{x}}|\mathbf{x}) \|\mathbf{s}_{\theta}(\tilde{\mathbf{x}})\|_2^2 d\tilde{\mathbf{x}} d\mathbf{x} \\ &= \frac{1}{2} E_{\mathbf{x} \sim p_{data}(\mathbf{x})} E_{\tilde{\mathbf{x}} \sim q_{\sigma}(\tilde{\mathbf{x}}|\mathbf{x})} [\|\mathbf{s}_{\theta}(\tilde{\mathbf{x}})\|_2^2]\end{aligned}$$

Denoising score matching

$$\begin{aligned} & - \int q_{\sigma}(\tilde{\mathbf{x}}) \nabla_{\tilde{\mathbf{x}}} \log q_{\sigma}(\tilde{\mathbf{x}})^T \mathbf{s}_{\theta}(\tilde{\mathbf{x}}) d\tilde{\mathbf{x}} = - \int \nabla_{\tilde{\mathbf{x}}} q_{\sigma}(\tilde{\mathbf{x}})^T \mathbf{s}_{\theta}(\tilde{\mathbf{x}}) d\tilde{\mathbf{x}} \\ & = - \int \nabla_{\tilde{\mathbf{x}}} \left(\int p_{data}(\mathbf{x}) q_{\sigma}(\tilde{\mathbf{x}}|\mathbf{x}) d\mathbf{x} \right)^T \mathbf{s}_{\theta}(\tilde{\mathbf{x}}) d\tilde{\mathbf{x}} \\ & = - \int \left(\int p_{data}(\mathbf{x}) \nabla_{\tilde{\mathbf{x}}} q_{\sigma}(\tilde{\mathbf{x}}|\mathbf{x}) d\mathbf{x} \right)^T \mathbf{s}_{\theta}(\tilde{\mathbf{x}}) d\tilde{\mathbf{x}} \\ & = - \int \left(\int p_{data}(\mathbf{x}) q_{\sigma}(\tilde{\mathbf{x}}|\mathbf{x}) \nabla_{\tilde{\mathbf{x}}} \log q_{\sigma}(\tilde{\mathbf{x}}|\mathbf{x}) d\mathbf{x} \right)^T \mathbf{s}_{\theta}(\tilde{\mathbf{x}}) d\tilde{\mathbf{x}} \\ & = - \int \int p_{data}(\mathbf{x}) q_{\sigma}(\tilde{\mathbf{x}}|\mathbf{x}) \nabla_{\tilde{\mathbf{x}}} \log q_{\sigma}(\tilde{\mathbf{x}}|\mathbf{x})^T \mathbf{s}_{\theta}(\tilde{\mathbf{x}}) d\tilde{\mathbf{x}} d\mathbf{x} \end{aligned}$$

Denoising score matching

$$\begin{aligned} & - \int q_{\sigma}(\tilde{\mathbf{x}}) \nabla_{\tilde{\mathbf{x}}} \log q_{\sigma}(\tilde{\mathbf{x}})^T \mathbf{s}_{\theta}(\tilde{\mathbf{x}}) d\tilde{\mathbf{x}} = \\ & = - \int \int p_{data}(\mathbf{x}) q_{\sigma}(\tilde{\mathbf{x}}|\mathbf{x}) \nabla_{\tilde{\mathbf{x}}} \log q_{\sigma}(\tilde{\mathbf{x}}|\mathbf{x})^T \mathbf{s}_{\theta}(\tilde{\mathbf{x}}) d\tilde{\mathbf{x}} d\mathbf{x} \\ & = - E_{\mathbf{x} \sim p_{data}(\mathbf{x})} E_{\tilde{\mathbf{x}} \sim q_{\sigma}(\tilde{\mathbf{x}}|\mathbf{x})} [\nabla_{\tilde{\mathbf{x}}} \log q_{\sigma}(\tilde{\mathbf{x}}|\mathbf{x})^T \mathbf{s}_{\theta}(\tilde{\mathbf{x}})] \end{aligned}$$

Denoising score matching

$$\begin{aligned} & \frac{1}{2} E_{\tilde{\mathbf{x}} \sim q_{\sigma}} [\|\nabla_{\tilde{\mathbf{x}}} \log q_{\sigma}(\tilde{\mathbf{x}}) - \mathbf{s}_{\theta}(\tilde{\mathbf{x}})\|_2^2] \\ &= \text{const.} + \frac{1}{2} E_{\mathbf{x} \sim p_{data}(\mathbf{x})} E_{\tilde{\mathbf{x}} \sim q_{\sigma}(\tilde{\mathbf{x}}|\mathbf{x})} [\|\mathbf{s}_{\theta}(\tilde{\mathbf{x}})\|_2^2] \\ &\quad - E_{\mathbf{x} \sim p_{data}(\mathbf{x})} E_{\tilde{\mathbf{x}} \sim q_{\sigma}(\tilde{\mathbf{x}}|\mathbf{x})} [\nabla_{\tilde{\mathbf{x}}} \log q_{\sigma}(\tilde{\mathbf{x}}|\mathbf{x})^T \mathbf{s}_{\theta}(\tilde{\mathbf{x}})] \\ &= \text{const.} + \frac{1}{2} E_{\mathbf{x} \sim p_{data}(\mathbf{x})} E_{\tilde{\mathbf{x}} \sim q_{\sigma}(\tilde{\mathbf{x}}|\mathbf{x})} [\|\nabla_{\tilde{\mathbf{x}}} \log q_{\sigma}(\tilde{\mathbf{x}}|\mathbf{x}) - \mathbf{s}_{\theta}(\tilde{\mathbf{x}})\|_2^2] \\ &\quad - \frac{1}{2} E_{\mathbf{x} \sim p_{data}(\mathbf{x})} E_{\tilde{\mathbf{x}} \sim q_{\sigma}(\tilde{\mathbf{x}}|\mathbf{x})} [\|\nabla_{\tilde{\mathbf{x}}} \log q_{\sigma}(\tilde{\mathbf{x}}|\mathbf{x})\|_2^2] \\ &= \frac{1}{2} E_{\mathbf{x} \sim p_{data}(\mathbf{x})} E_{\tilde{\mathbf{x}} \sim q_{\sigma}(\tilde{\mathbf{x}}|\mathbf{x})} [\|\nabla_{\tilde{\mathbf{x}}} \log q_{\sigma}(\tilde{\mathbf{x}}|\mathbf{x}) - \mathbf{s}_{\theta}(\tilde{\mathbf{x}})\|_2^2] + \text{const.} \end{aligned}$$

Denoising score matching

- Estimate the score of a noise-perturbed distribution

$$\min_{\theta} E_{\tilde{\mathbf{x}} \sim q_{\sigma}} [\|\nabla_{\tilde{\mathbf{x}}} \log q_{\sigma}(\tilde{\mathbf{x}}) - \mathbf{s}_{\theta}(\tilde{\mathbf{x}})\|_2^2]$$
$$\Leftrightarrow \min_{\theta} E_{\mathbf{x} \sim p_{data}(\mathbf{x})} E_{\tilde{\mathbf{x}} \sim q_{\sigma}(\tilde{\mathbf{x}}|\mathbf{x})} [\|\nabla_{\tilde{\mathbf{x}}} \log q_{\sigma}(\tilde{\mathbf{x}}|\mathbf{x}) - \mathbf{s}_{\theta}(\tilde{\mathbf{x}})\|_2^2]$$

- $\nabla_{\tilde{\mathbf{x}}} \log q_{\sigma}(\tilde{\mathbf{x}}|\mathbf{x})$ is easy to compute

$$q_{\sigma}(\tilde{\mathbf{x}}|\mathbf{x}) = N(\tilde{\mathbf{x}}|\mathbf{x}, \sigma^2 \mathbf{I}), \quad \nabla_{\tilde{\mathbf{x}}} \log q_{\sigma}(\tilde{\mathbf{x}}|\mathbf{x}) = -\frac{1}{\sigma^2} (\tilde{\mathbf{x}} - \mathbf{x})$$

- **Pros:** efficient to optimize even for very high dimensional data, and useful for optimal denoising.
- **Con:** cannot estimate the score of clean data (noise-free)

Training for denoising score matching

- Sample a minibatch of datapoints $\mathbf{x}^{(1)}, \mathbf{x}^{(2)}, \dots, \mathbf{x}^{(n)}$ from p_{data}
- Sample a minibatch of perturbed datapoints $\tilde{\mathbf{x}}^{(1)}, \tilde{\mathbf{x}}^{(2)}, \dots, \tilde{\mathbf{x}}^{(n)}$ from $q_{\sigma}(\tilde{\mathbf{x}})$

$$\tilde{\mathbf{x}}^{(i)} \sim q_{\sigma}(\tilde{\mathbf{x}}^{(i)} | \mathbf{x}^{(i)}) = N(\tilde{\mathbf{x}}^{(i)} | \mathbf{x}^{(i)}, \sigma^2 \mathbf{I})$$

- Estimate the denoising score matching loss with empirical means

$$\frac{1}{2n} \sum_{i=1}^n \left\| \nabla_{\tilde{\mathbf{x}}} \log q_{\sigma}(\tilde{\mathbf{x}}^{(i)} | \mathbf{x}^{(i)}) - \mathbf{s}_{\theta}(\tilde{\mathbf{x}}^{(i)}) \right\|$$

- I.e.,

$$\frac{1}{2n} \sum_{i=1}^n \left\| \frac{1}{\sigma^2} (\tilde{\mathbf{x}}^{(i)} - \mathbf{x}^{(i)}) + \mathbf{s}_{\theta}(\tilde{\mathbf{x}}^{(i)}) \right\|$$

- Stochastic gradient descent
- Need to choose a small σ

Pitfall of denoising score matching

- The loss variance will increase drastically as $\sigma \rightarrow 0$
- Denoising score matching loss for Gaussian perturbations

$$\begin{aligned} & E_{\mathbf{x} \sim p_{data}(\mathbf{x})} E_{\tilde{\mathbf{x}} \sim q_{\sigma}(\tilde{\mathbf{x}}|\mathbf{x})} \left[\left\| \frac{1}{\sigma^2} (\tilde{\mathbf{x}} - \mathbf{x}) + \mathbf{s}_{\theta}(\tilde{\mathbf{x}}) \right\|_2^2 \right] \\ &= E_{\mathbf{x} \sim p_{data}(\mathbf{x})} E_{\mathbf{z} \sim N(\mathbf{0}, I)} \left[\left\| \frac{1}{\sigma} \mathbf{z} + \mathbf{s}_{\theta}(\mathbf{x} + \sigma \mathbf{z}) \right\|_2^2 \right] \\ &= E_{\mathbf{x} \sim p_{data}(\mathbf{x})} E_{\mathbf{z} \sim N(\mathbf{0}, I)} \left[\|\mathbf{s}_{\theta}(\mathbf{x} + \sigma \mathbf{z})\|_2^2 + 2\mathbf{s}_{\theta}(\mathbf{x} + \sigma \mathbf{z})^T \frac{\mathbf{z}}{\sigma} + \frac{\|\mathbf{z}\|_2^2}{\sigma^2} \right] \end{aligned}$$

- If we choose very small $\sigma \rightarrow 0$

$$\text{Var} \left(\frac{\mathbf{z}}{\sigma} \right) \rightarrow \infty,$$

$$E_{\mathbf{z} \sim N(\mathbf{0}, I)} \left[\frac{\|\mathbf{z}\|_2^2}{\sigma^2} \right] \rightarrow \infty$$

Denoising score matching

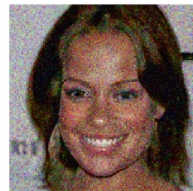
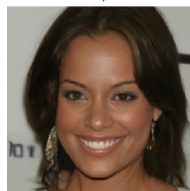
- Consider the perturbed distribution

$$q_{\sigma}(\tilde{\mathbf{x}}|\mathbf{x}) := N(\tilde{\mathbf{x}}|\mathbf{x}, \sigma^2 \mathbf{I}), \quad q_{\sigma}(\tilde{\mathbf{x}}) = \int p_{data}(\mathbf{x}) q_{\sigma}(\tilde{\mathbf{x}}|\mathbf{x}) d\mathbf{x}$$

- Score estimation for $\nabla_{\tilde{\mathbf{x}}} \log q_{\sigma}(\tilde{\mathbf{x}})$
- Score matching

$$\min_{\theta} E_{\tilde{\mathbf{x}} \sim q_{\sigma}} [\|\nabla_{\tilde{\mathbf{x}}} \log q_{\sigma}(\tilde{\mathbf{x}}) - \mathbf{s}_{\theta}(\tilde{\mathbf{x}})\|_2^2]$$

$$\Leftrightarrow \min_{\theta} E_{\mathbf{x} \sim p_{data}(\mathbf{x})} E_{\tilde{\mathbf{x}} \sim q_{\sigma}(\tilde{\mathbf{x}}|\mathbf{x})} \left[\left\| \frac{1}{\sigma^2} (\tilde{\mathbf{x}} - \mathbf{x}) + \mathbf{s}_{\theta}(\tilde{\mathbf{x}}) \right\|_2^2 \right]$$



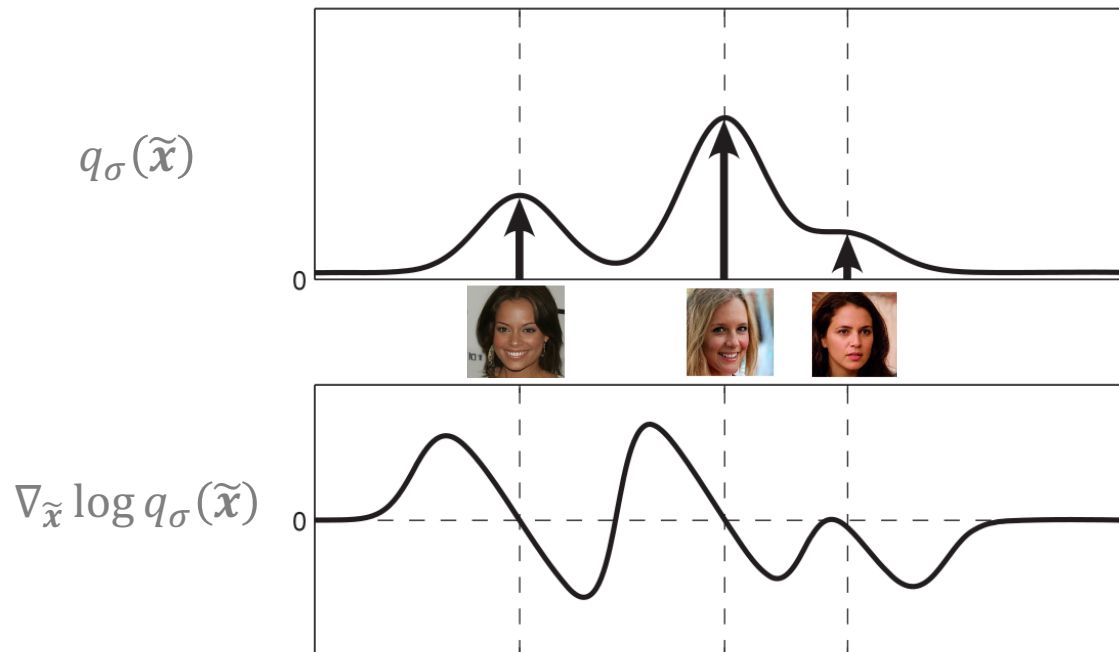
- $\mathbf{s}_{\theta}(\tilde{\mathbf{x}})$ tries to estimate the noise that was added to produce $\tilde{\mathbf{x}}$

Tweedie's formula

- Score estimation for $\nabla_{\tilde{\mathbf{x}}} \log q_{\sigma}(\tilde{\mathbf{x}})$ is equivalent to denoising

$$E_{\mathbf{x} \sim p_{data}(\mathbf{x})} E_{\tilde{\mathbf{x}} \sim q_{\sigma}(\tilde{\mathbf{x}}|\mathbf{x})} \left[\left\| \frac{1}{\sigma^2} (\tilde{\mathbf{x}} - \mathbf{x}) + \mathbf{s}_{\theta}(\tilde{\mathbf{x}}) \right\|_2^2 \right]$$

- Tweedie's formula:** Optimal denoising strategy is to follow the gradient (score):



Tweedie's formula and denoising score matching

- Denoising score matching is suitable for optimal denoising
- Given $p_{data}(\mathbf{x})$ and $q_{\sigma}(\tilde{\mathbf{x}}|\mathbf{x})$, we can define the posterior $p(\mathbf{x}|\tilde{\mathbf{x}})$ with Bayes' rule

$$p(\mathbf{x}|\tilde{\mathbf{x}}) \propto p_{data}(\mathbf{x})q_{\sigma}(\tilde{\mathbf{x}}|\mathbf{x})$$

- Recall that

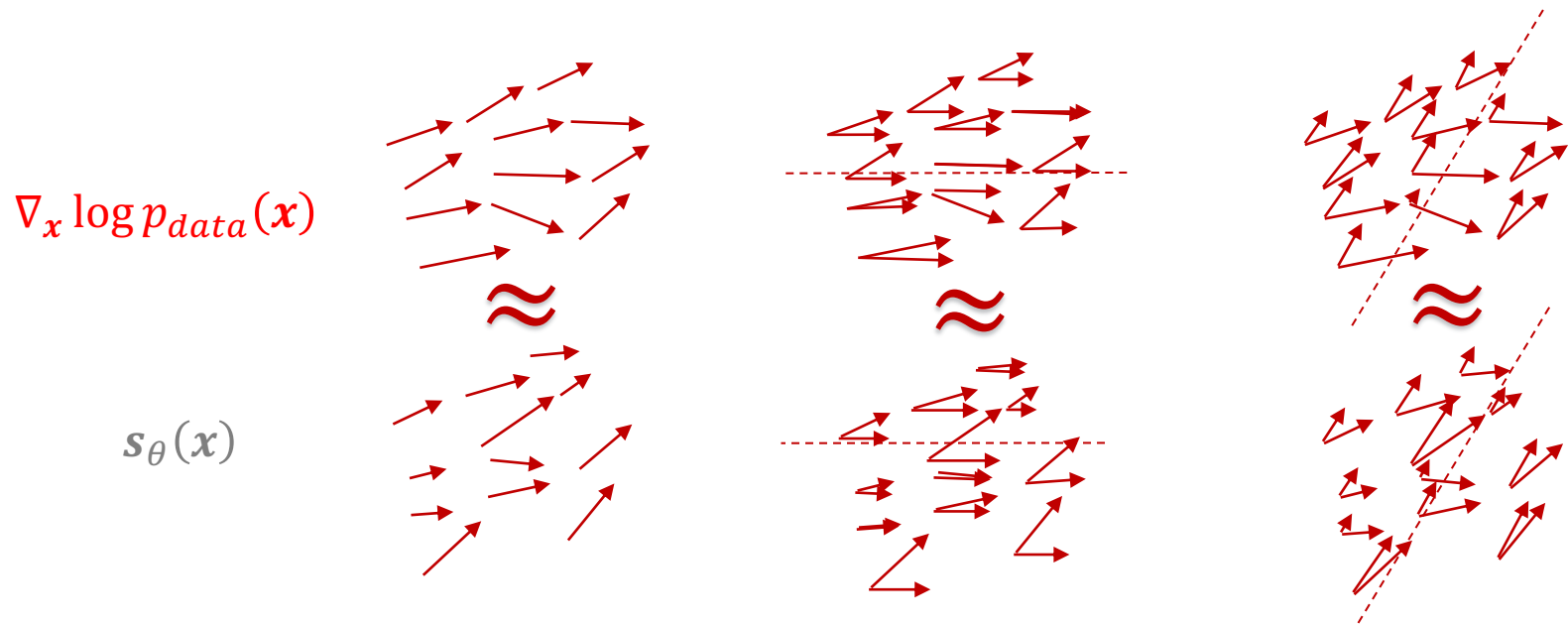
$$q_{\sigma}(\tilde{\mathbf{x}}) = \int p_{data}(\mathbf{x})q_{\sigma}(\tilde{\mathbf{x}}|\mathbf{x})d\mathbf{x}$$

- Tweedie's formula:

$$E_{\mathbf{x} \sim p(\mathbf{x}|\tilde{\mathbf{x}})}[\mathbf{x}] = \tilde{\mathbf{x}} + \sigma^2 \nabla_{\mathbf{x}} \log q_{\sigma}(\tilde{\mathbf{x}}) \approx \tilde{\mathbf{x}} + \sigma^2 s_{\theta}(\tilde{\mathbf{x}})$$

Sliced score matching

- One dimensional problems is easier
- Consider projections onto random directions



Song*, Garg*, Shi, Ermon. "Sliced Score Matching: A Scalable Approach to Density and Score Estimation." UAI 2019.

Sliced score matching

- **Objective:** Sliced Fisher Divergence

$$\frac{1}{2} E_{v \sim p_v} E_{x \sim p_{data}} \left[\left(v^T \nabla_x \log p_{data}(x) - v^T s_{\theta}(x) \right)^2 \right]$$

- Integration by parts

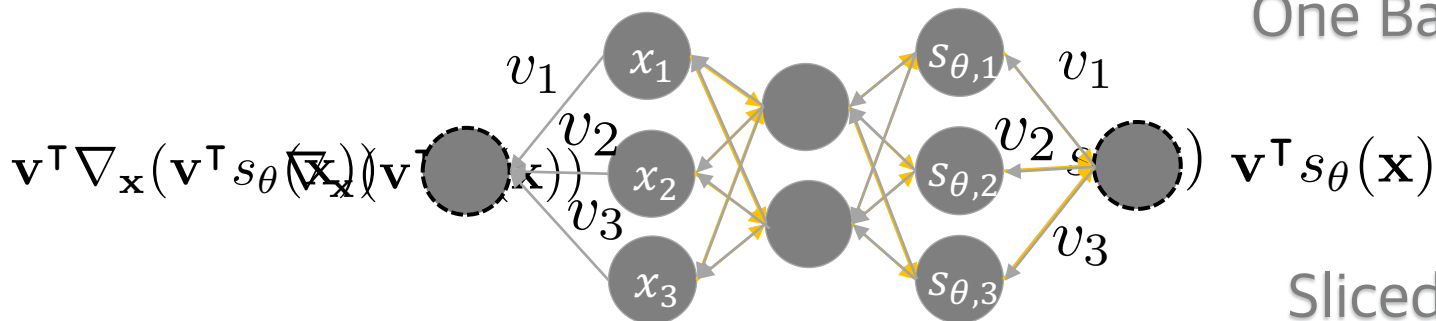
$$E_{v \sim p_v} E_{x \sim p_{data}} \left[\frac{1}{2} \left(v^T s_{\theta}(x) \right)^2 + v^T \nabla_x s_{\theta}(x) v \right]$$

$$(v_1 \quad v_2 \quad v_3) \begin{pmatrix} \frac{\partial s_{\theta,1}(x)}{\partial x_1} & \frac{\partial s_{\theta,1}(x)}{\partial x_2} & \frac{\partial s_{\theta,1}(x)}{\partial x_3} \\ \frac{\partial s_{\theta,2}(x)}{\partial x_1} & \frac{\partial s_{\theta,2}(x)}{\partial x_2} & \frac{\partial s_{\theta,2}(x)}{\partial x_3} \\ \frac{\partial s_{\theta,3}(x)}{\partial x_1} & \frac{\partial s_{\theta,3}(x)}{\partial x_2} & \frac{\partial s_{\theta,3}(x)}{\partial x_3} \end{pmatrix} \begin{pmatrix} v_1 \\ v_2 \\ v_3 \end{pmatrix}$$

Computing Jacobian-vector products is scalable

$$\mathbf{v}^T \nabla_{\mathbf{x}} \mathbf{s}_{\theta}(\mathbf{x}) \mathbf{v} = \mathbf{v}^T \nabla_{\mathbf{x}} (\mathbf{v}^T \mathbf{s}_{\theta}(\mathbf{x}))$$

One Backprop!

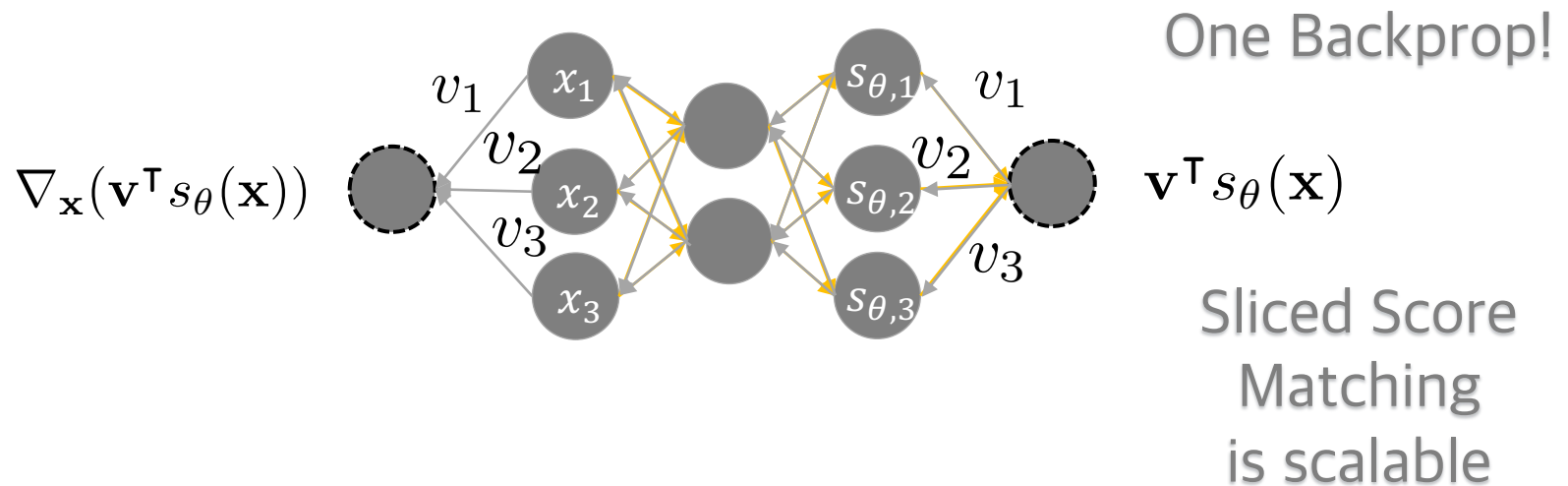


Sliced Score
Matching
is scalable

- Slightly slower than denoising score matching

Computing Jacobian-vector products is scalable

$$\mathbf{v}^T \nabla_{\mathbf{x}} \mathbf{s}_{\theta}(\mathbf{x}) \mathbf{v} = \mathbf{v}^T \nabla_{\mathbf{x}} (\mathbf{v}^T \mathbf{s}_{\theta}(\mathbf{x}))$$



- Slightly slower than denoising score matching

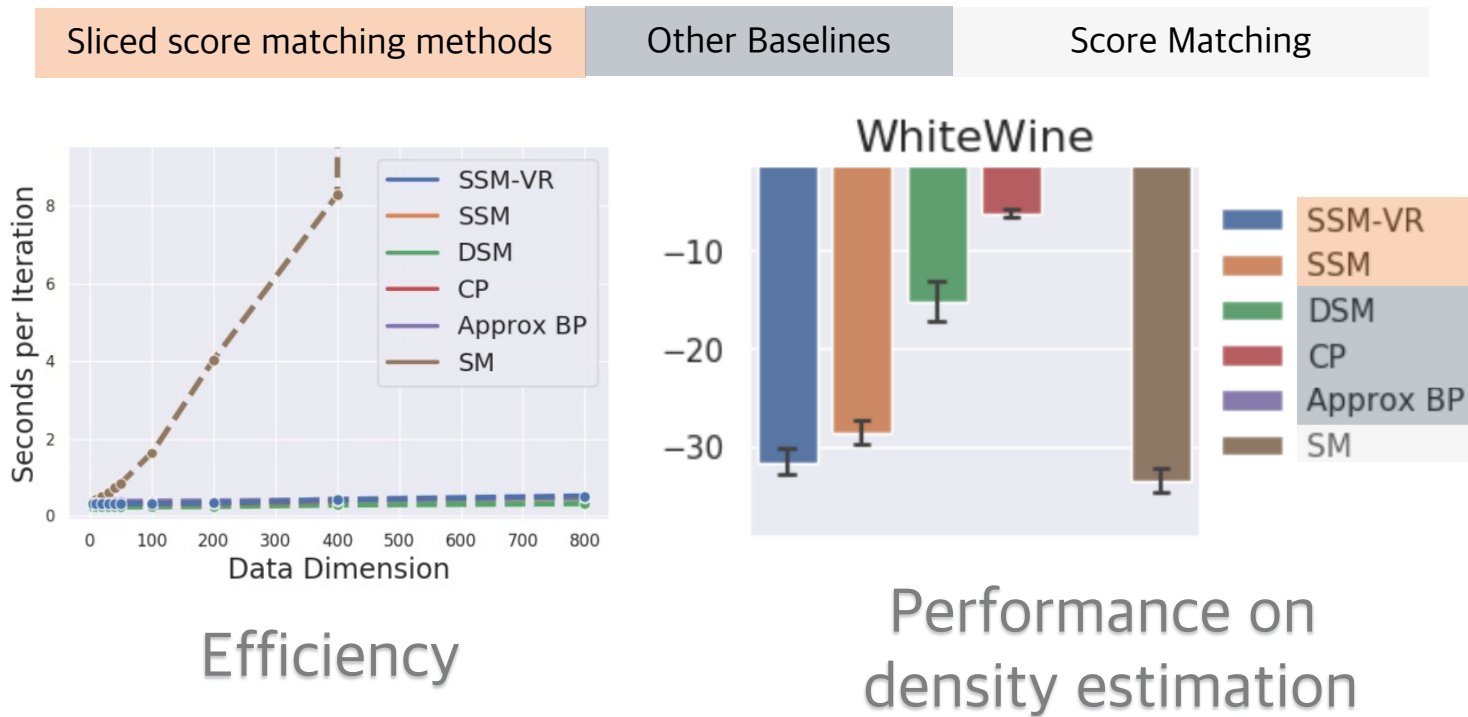
Training of sliced score matching

- Sample a minibatch of datapoints $\mathbf{x}^{(1)}, \mathbf{x}^{(2)}, \dots, \mathbf{x}^{(n)}$ from p_{data}
- Sample a minibatch of projection directions $\mathbf{v}^{(1)}, \mathbf{v}^{(2)}, \dots, \mathbf{v}^{(n)}$ from p_v
- Estimate the sliced score matching loss with empirical means

$$\frac{1}{n} \sum_{i=1}^n \mathbf{v}^{(i)T} \nabla_{\mathbf{x}} \mathbf{s}_{\theta}(\mathbf{x}^{(i)}) \mathbf{v}^{(i)} + \frac{1}{2} \left(\mathbf{v}^{(i)T} \mathbf{s}_{\theta}(\mathbf{x}^{(i)}) \right)^2$$

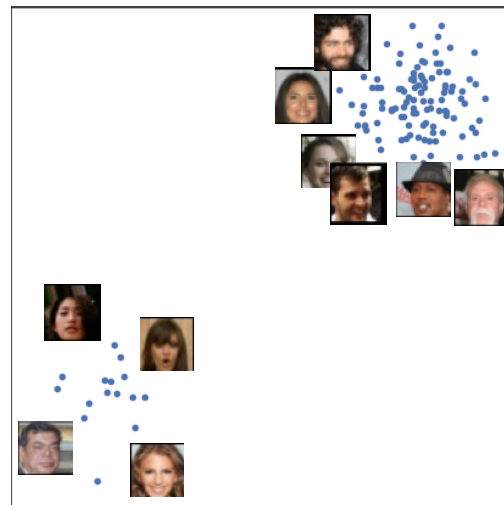
- The projection distribution is typically Gaussian or Rademacher
- Stochastic gradient descent
- Can use more projections per datapoint to boost performance

Experimental results for sliced score matching



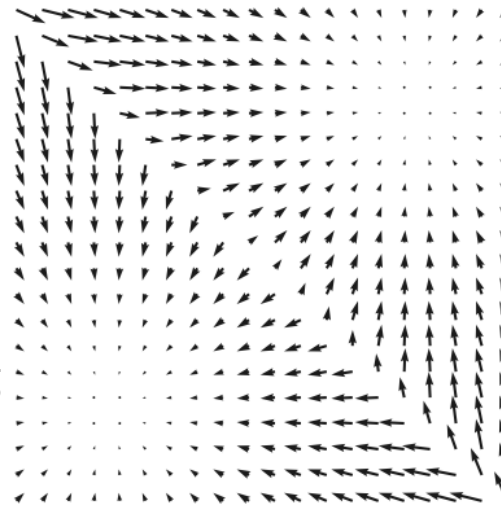
Song*, Garg*, Shi, Ermon. "Sliced Score Matching: A Scalable Approach to Density and Score Estimation." UAI 2019.

Score-based generative modeling



Data samples
 $\mathbf{x}^{(1)}, \mathbf{x}^{(2)}, \dots, \mathbf{x}^{(N)}$

Score
Matching



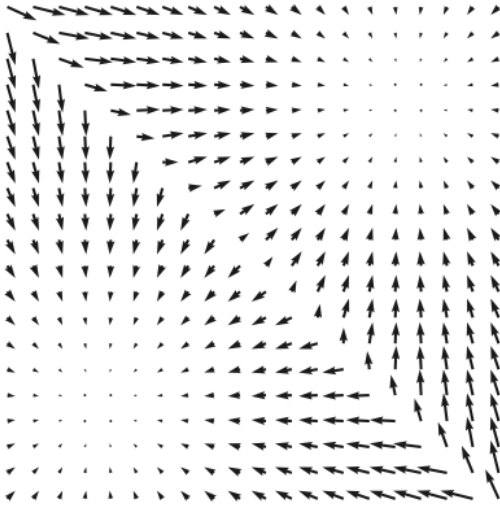
Scores function
 $\mathbf{s}_\theta(\mathbf{x}) \approx \nabla_{\mathbf{x}} \log p_{data}(\mathbf{x})$

?

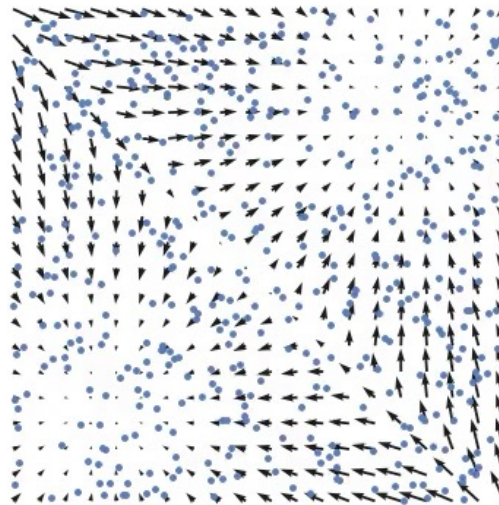


New samples

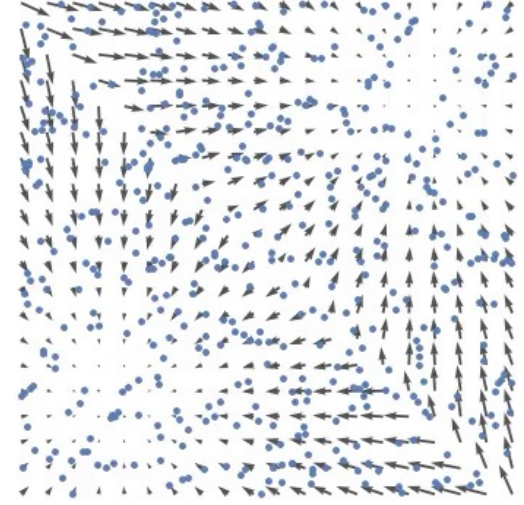
From scores to samples: Langevin MCMC



Scores
 $\mathbf{s}_\theta(\mathbf{x})$



Follow the scores
 $\mathbf{x}^{t+1} = \mathbf{x}^t + \frac{\epsilon}{2} \mathbf{s}_\theta(\mathbf{x})$

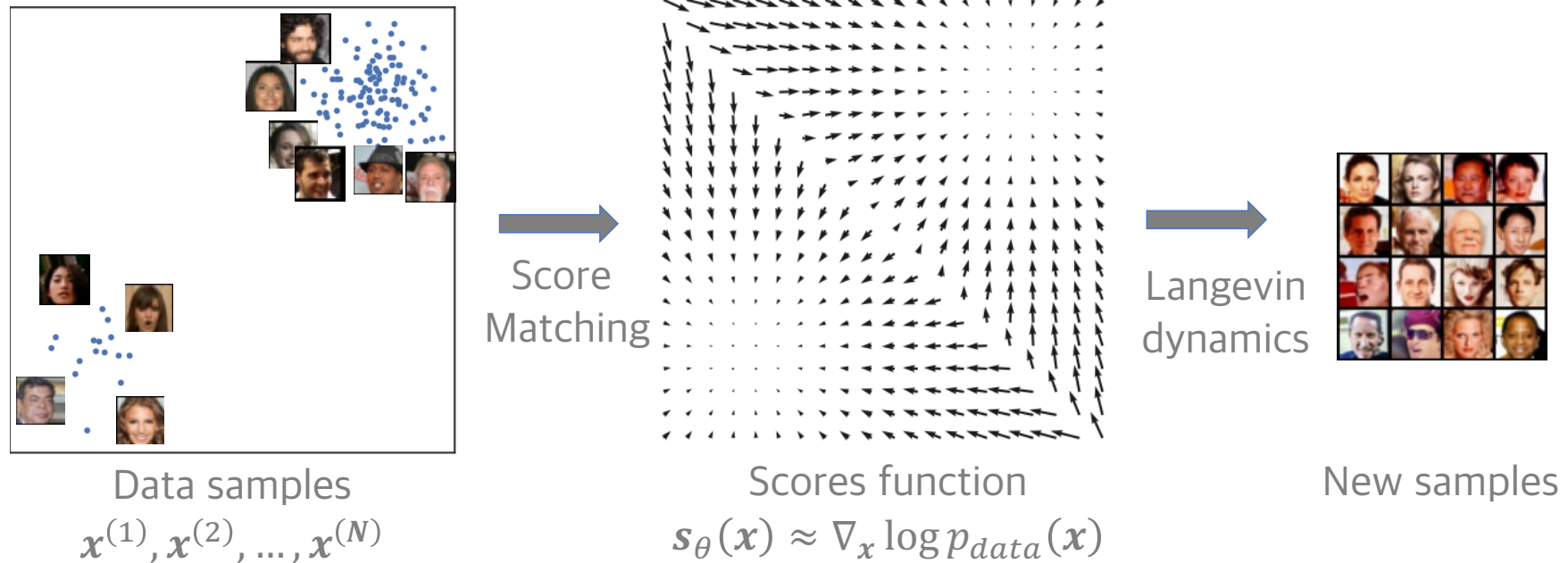


Follow noisy scores:
Langevin MCMC
 $\mathbf{z} \sim N(\mathbf{0}, I)$
 $\mathbf{x}^{t+1} = \mathbf{x}^t + \frac{\epsilon}{2} \mathbf{s}_\theta(\mathbf{x}) + \sqrt{\epsilon} \mathbf{z}$

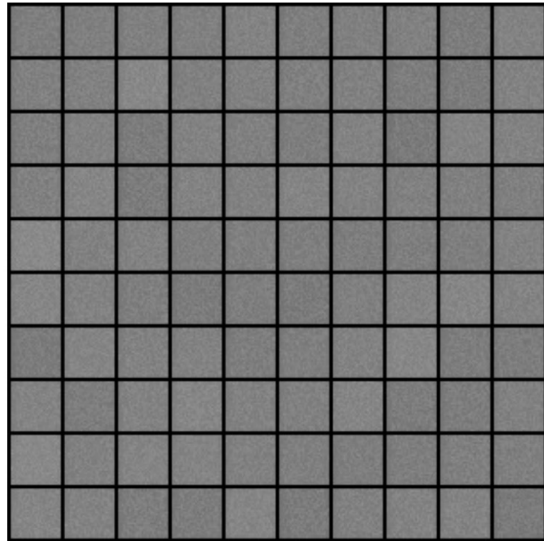
Langevin dynamics sampling

- Sample from $p_{data}(\mathbf{x})$ using only the scores $\nabla_{\mathbf{x}} \log p_{data}(\mathbf{x})$
 - Initialize $\mathbf{x}^0 \sim \pi(\mathbf{x})$
 - Repeat for $t = 0, \dots, T - 1$
 - $\mathbf{z} \sim N(0, I)$
 - $\mathbf{x}^{t+1} = \mathbf{x}^t + \frac{\epsilon}{2} \nabla_{\mathbf{x}} \log p_{data}(\mathbf{x})|_{\mathbf{x}=\mathbf{x}^t} + \sqrt{\epsilon} \mathbf{z}$
 - If $\epsilon \rightarrow 0$ and $T \rightarrow \infty$, then we have $\mathbf{x}^T$ converges to a sample from p_{data}
- Langevin dynamics + score estimation
$$\mathbf{s}_{\theta}(\mathbf{x}) \approx \nabla_{\mathbf{x}} \log p_{data}(\mathbf{x})$$

Score-based generative modeling



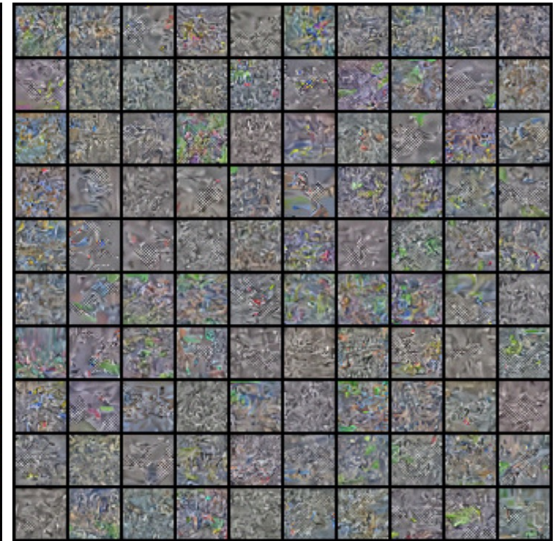
Score-based generative modeling: results



(a) MNIST



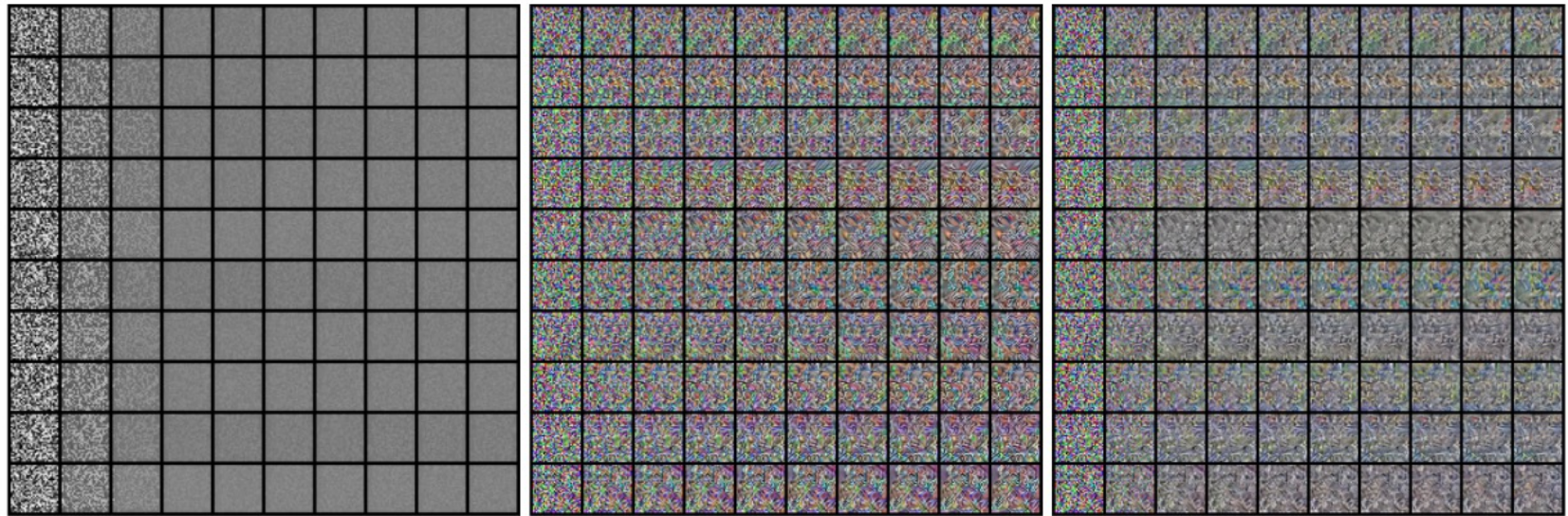
(b) CelebA



(c) CIFAR-10

Final samples

Score-based generative modeling: results



(a) MNIST

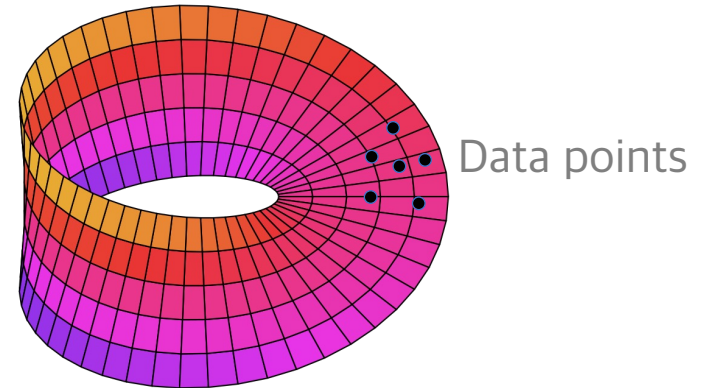
(b) CelebA

(c) CIFAR-10

Langevin sampling process

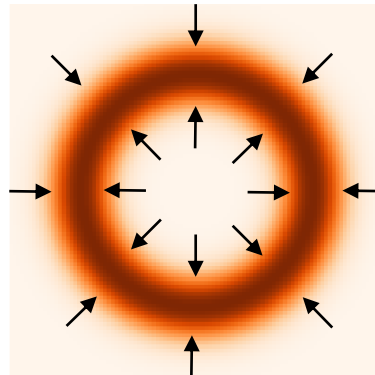
Pitfall 1: manifold hypothesis

- Manifold hypothesis



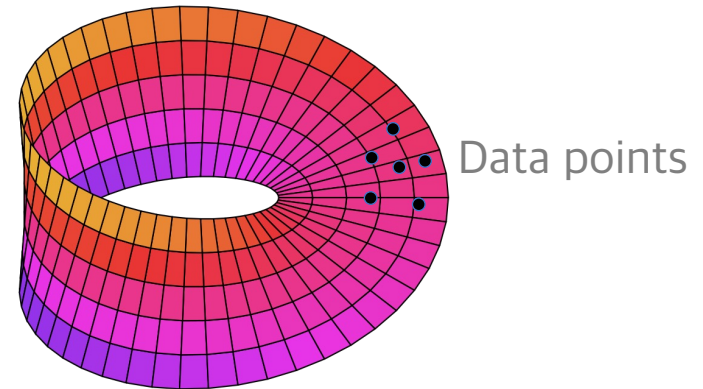
- Data score is undefined

$$\nabla_x \log p_{data}(x)$$



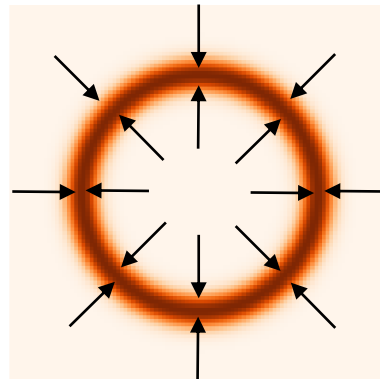
Pitfall 1: manifold hypothesis

- Manifold hypothesis



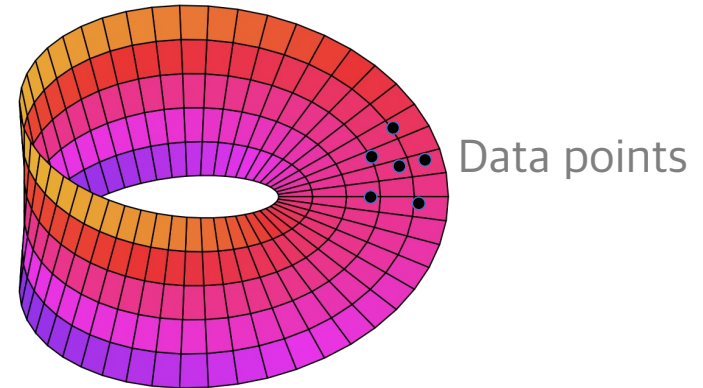
- Data score is undefined

$$\nabla_x \log p_{data}(x)$$

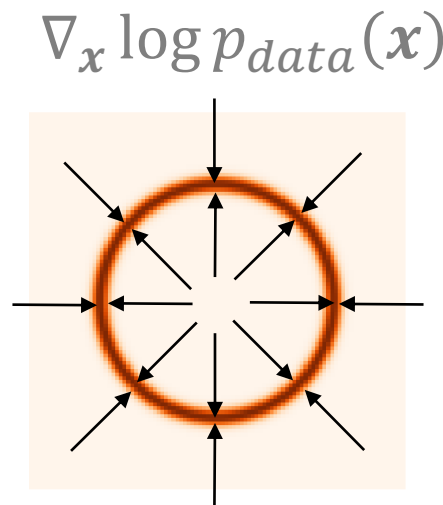


Pitfall 1: manifold hypothesis

- Manifold hypothesis



- Data score is undefined

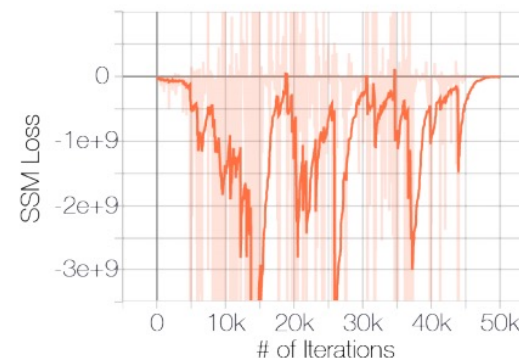


Pitfall 1: manifold hypothesis

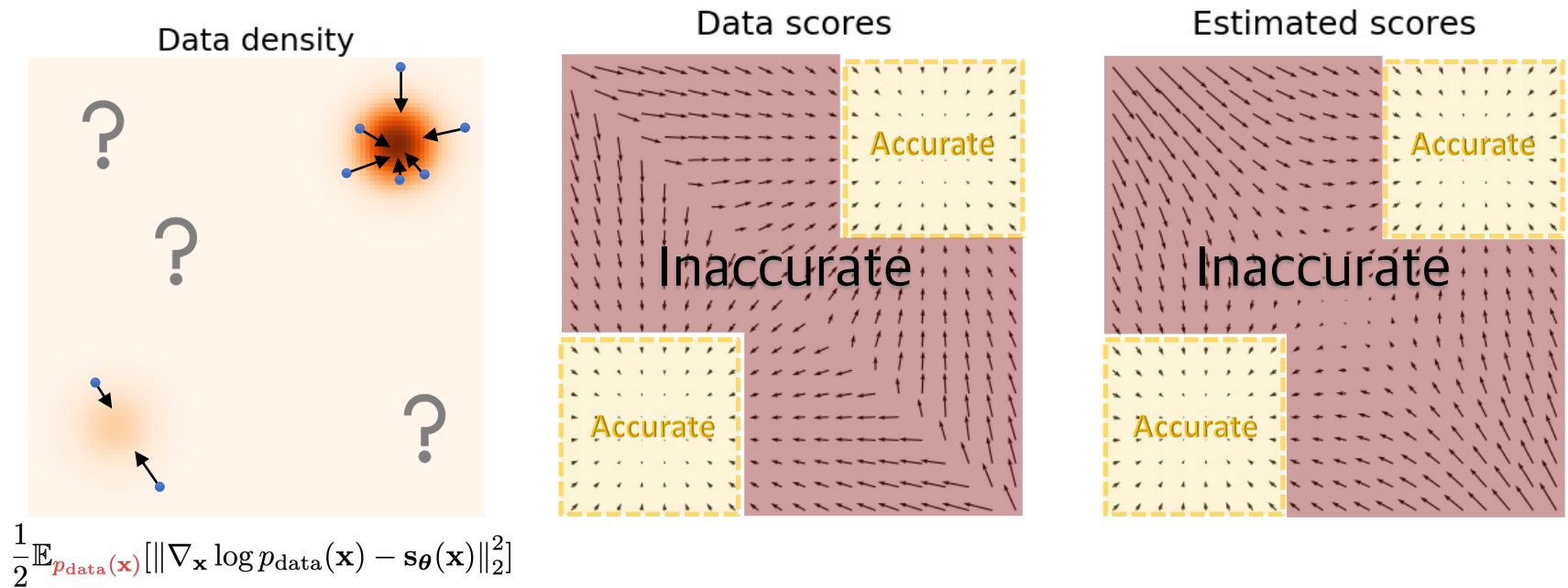
- Fitting the data with a low-dimensional linear manifold (PCA)



- Score estimation on CIFAR-10



Pitfall 2: challenge in low data density regions



Langevin MCMC will have trouble
exploring low density regions

Song and Ermon. "Generative Modeling by Estimating Gradients of the Data Distribution." NeurIPS 2019.

Pitfall 3: slow mixing of Langevin dynamics between data modes

- Suppose the data distribution has two modes with disjoint supports:

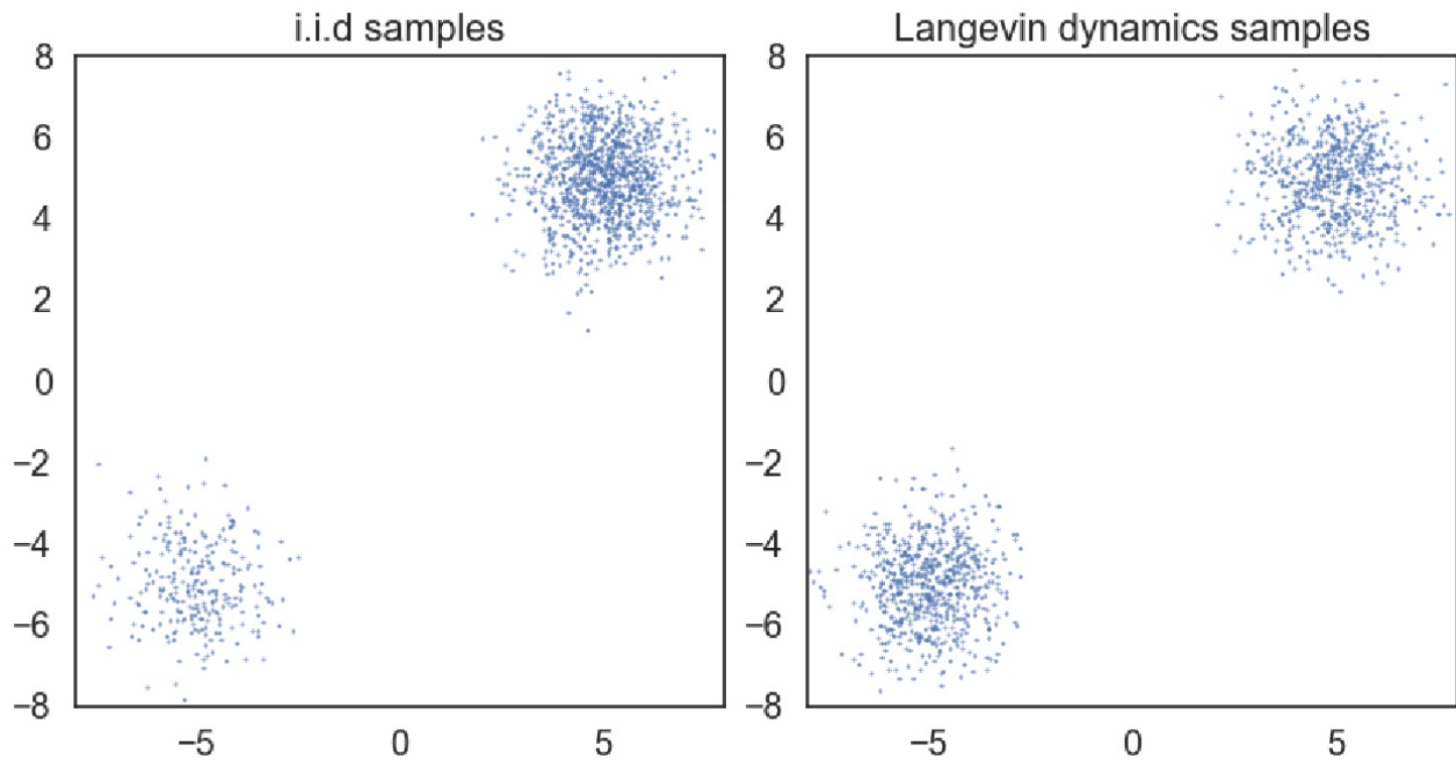
$$p_{data}(\mathbf{x}) = \pi p_1(\mathbf{x}) + (1 - \pi)p_2(\mathbf{x})$$
$$A \cap B = \emptyset, \quad p_{data}(x) = \begin{cases} \pi p_1(x) & x \in A \\ (1 - \pi)p_2(x) & x \in B \end{cases}$$

- Data score function:

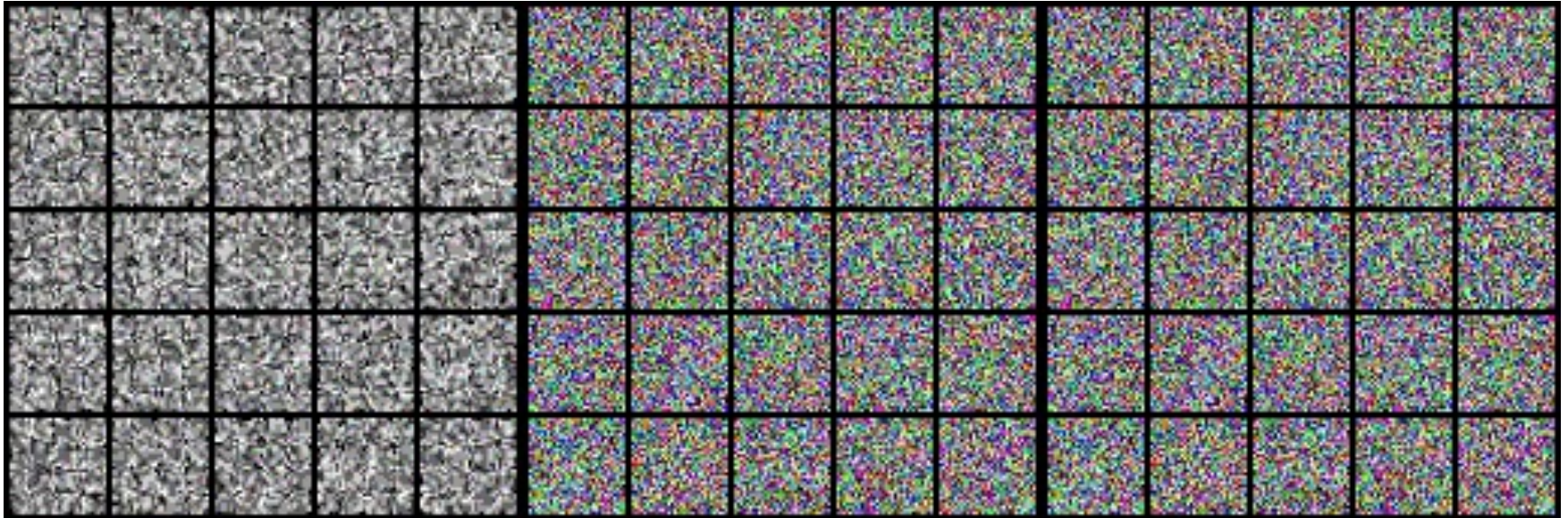
$$\nabla_x \log p_{data}(\mathbf{x}) = \begin{cases} \nabla_x \log p_1(\mathbf{x}) & x \in A \\ \nabla_x \log p_2(\mathbf{x}) & x \in B \end{cases}$$

- The score function has no dependence on π
- Langevin sampling will not reflect π

Pitfall 3: slow mixing of Langevin dynamics between data modes



After fixing pitfalls



Song, Yang, and Stefano Ermon. "Generative Modeling by Estimating Gradients of the Data Distribution." NeurIPS 2019

Thanks
