



CSCI- B 657 - Discussion Section

Vibhas Vats

Drawback of VAEs

- Objective function of VAEs

The diagram illustrates the components of the VAE objective function. At the top, 'Encoder: LV given input image x_i ' has a red arrow pointing to the latent variable $\mathbf{z}$ in the KL-divergence term of the equation below. To the right, $N(0,1)$ has a red arrow pointing to the same $\mathbf{z}$ and another red arrow pointing to the word 'prior'. The word 'prior' is positioned below $N(0,1)$. The equation is
$$\mathcal{L}_i(\theta, \phi) = -\mathbb{E}_{\mathbf{z} \sim q_\theta(\mathbf{z}|\mathbf{x}_i)}[\log p_\phi(\mathbf{x}_i|\mathbf{z})] + \underbrace{KL(q_\theta(\mathbf{z}|\mathbf{x}_i)||p(\mathbf{z}))}$$

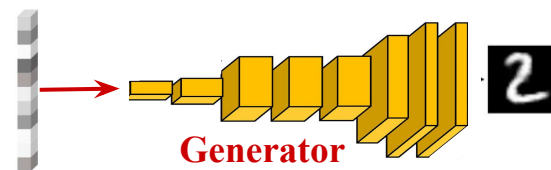
- If KL-div is large $\Rightarrow$ no proper optimization $\Rightarrow$ poor image generation

Why?

- very strong prior distribution $p(\mathbf{z}) \rightarrow$ prior distribution *different* than the data distribution
 - if the prior is **not** proper $\rightarrow$ poor optimization $\rightarrow$ poor reconstruction
- another probabilistic explanation exists
 - involving evidence lower bound (ELBO) ❌
- GANs can produce better results

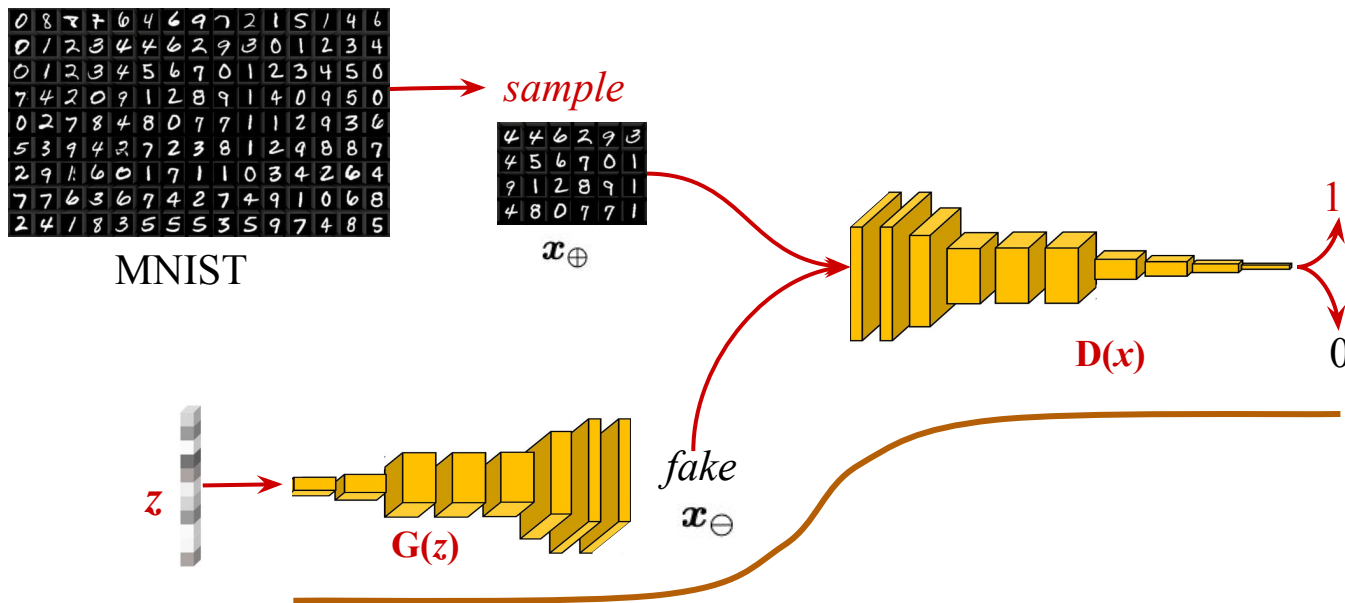
How GAN does it?

- **Generator (G):**
 - from an N -dimensional code $\rightarrow$ generate images
 - *main goal*: to produce images that looks real
 - real to whom? $\rightarrow$ **to human perception**
- But we can't have human in a training loop!
 - Another network $\rightarrow$ acts as *proxy* for human
 - **aim**: to discriminate b/w *real* and *fake* images
- **Discriminator (D):**
 - How to train such a network?
 - *input*:
 - a set of *real* images $\rightarrow$ from data distribution
 - and a set of *fake* images $\rightarrow$ output of G
 - ask to perform a *binary classification*:
 - real $\rightarrow$ 1, fake $\rightarrow$ 0
 - Sigmoid is used for binary classification



Competition b/w G and D

- **Goal of D:**→ to correctly classify real and fake images
 - **Goal of G:**→ to fool the D
 - fool the D to classify a fake as real
 - This has to be done simultaneously!
- If G can fool D
 - and we know that D is good enough
 - then, G understand the data distribution



0	8	7	6	4	6	9	7	2	1	5	1	4	6
0	1	2	3	4	6	2	9	3	0	1	2	3	4
0	1	2	3	4	5	6	7	0	1	2	3	4	5
7	4	2	0	9	1	2	8	9	1	4	0	9	5
0	2	7	8	4	8	0	7	7	1	1	2	9	3
5	3	9	4	2	7	2	3	8	1	2	9	8	7
2	9	1	6	0	1	7	1	1	0	3	4	2	6
7	7	6	3	6	7	4	2	7	4	9	1	0	6
2	4	1	8	3	5	5	3	5	9	7	4	8	5

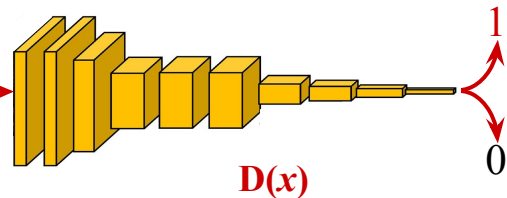
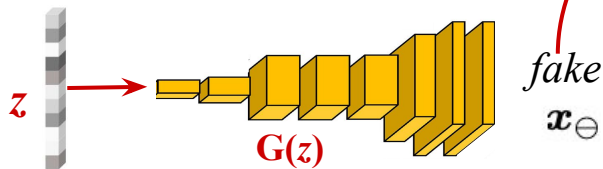
MNIST

sample

4	4	6	2	9	3
4	5	6	7	0	1
9	1	2	8	9	1
4	8	0	7	7	1

$x_{\oplus}$

Competition b/w G and D

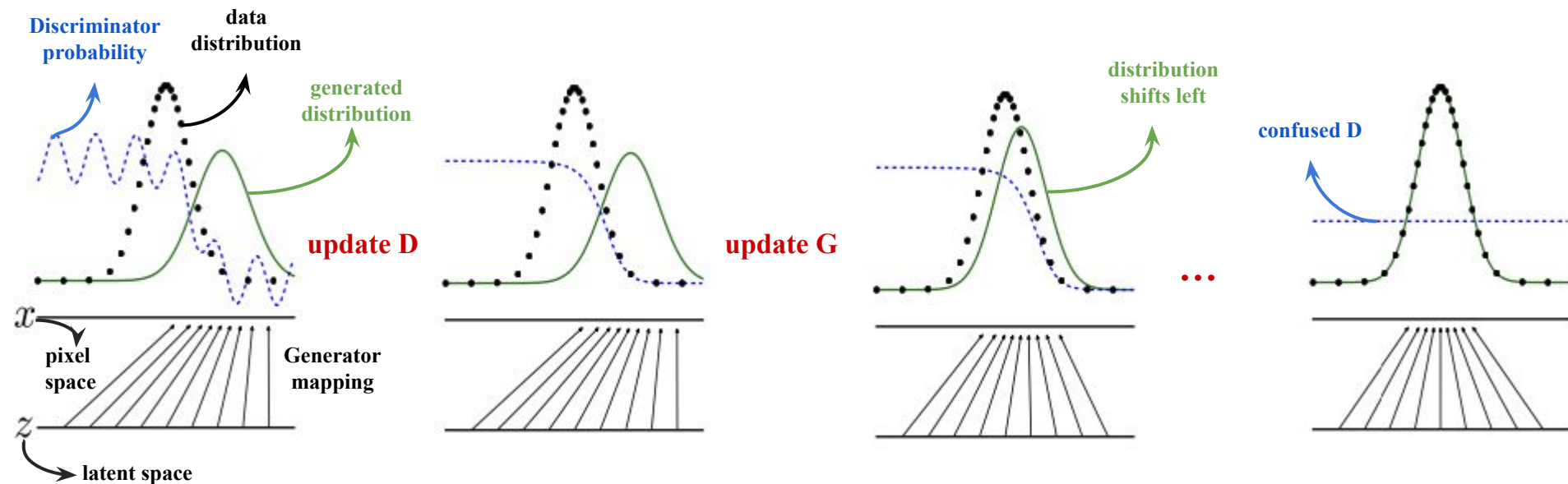


- If the generator does a good job
 - creates fake examples that looks like real
 - harms discriminator: $D(x_{-}) \uparrow$
- Dilemma:
 - discriminator needs the generator to be good $\rightarrow$ to be able to learn from confusing fake examples
 - but, by definition a good generator promotes $D(x_{-}) \uparrow$
 - harms the performance of the discriminator $\Rightarrow$ 2-player minimax game
- If the discriminator does a good job
 - becomes good at discerning x_{+} and x_{-}
 - $D(x_{+}) \uparrow$ and,
 - harms generator: $D(x_{-}) \downarrow$

GAN Optimization

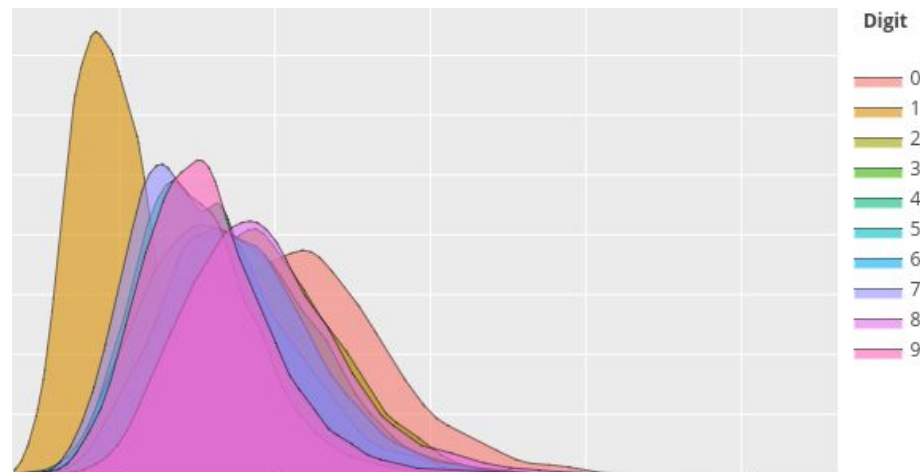
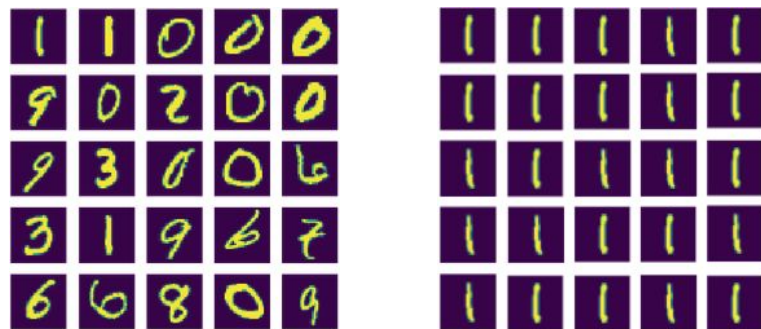
- Aim of GAN optimization is to achieve **Nash equilibrium**:
 - $D(x+) = D(x-) = 0.5$
 - discriminator can't tell real and fake apart!!
 - *meaning*: generator can generate images similar to real ones

Nobel Prize 1994:
“for pioneering analysis of
equilibria in the theory of
non-cooperative games”



Mode Collapse

- Which set of generator images do you prefer?



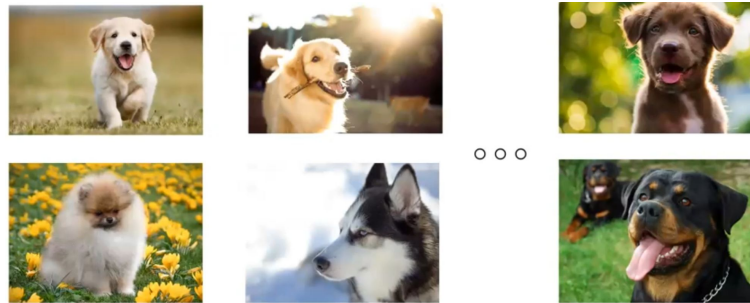
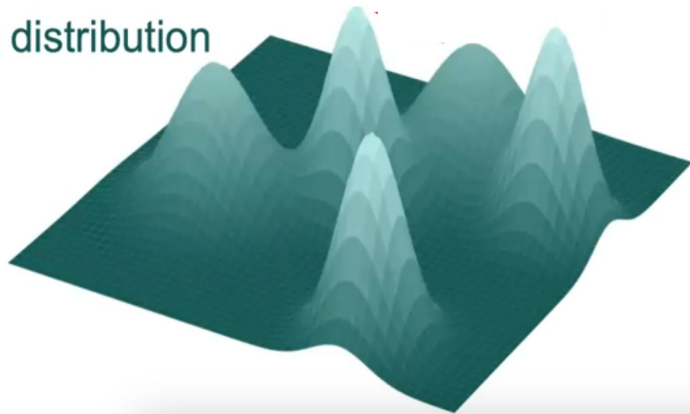
- A too much optimized generator can get stuck at a model of the generative model
 - it will generate real-looking fake images but mostly single class
 - initially D will classify it 'real'
 - but then it realizes and catches it up, but the G can jump to another **mode** ...
- GAN generates sharp realistic images
 - but these images do not represent the *true data distribution*
- Very careful optimization is needed to handle this

Estimating Probability Distribution of Data

Data distribution
(unknown)



Model distribution



Dataset

Generate

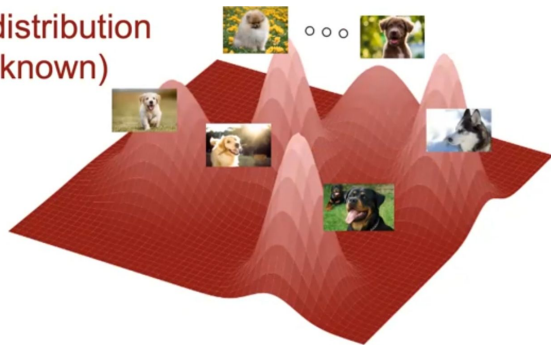


Novel data points

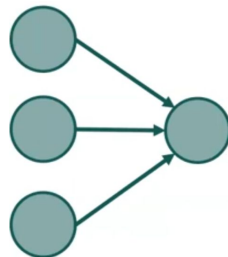
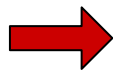
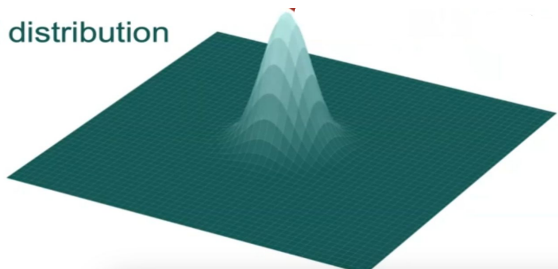


Complex Generative Models

Data distribution
(unknown)

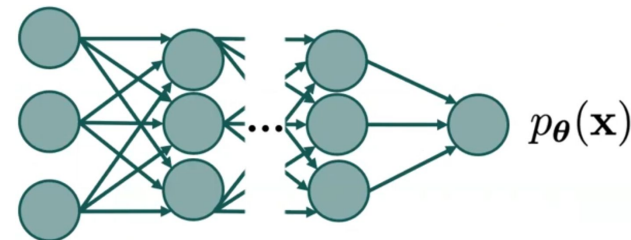
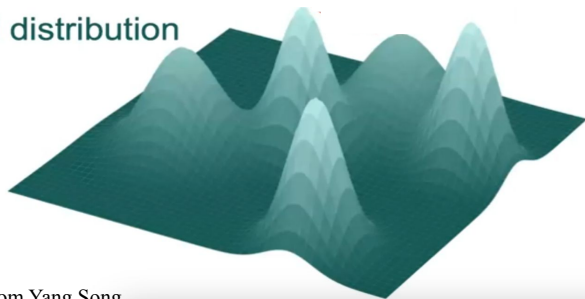


Model distribution



$$p_{\mu}(\mathbf{x}) = \frac{1}{(2\pi)^{d/2}} e^{-\frac{\|\mathbf{x}-\mu\|_2^2}{2}}$$

Model distribution



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

Deep Neural Network

- Data Distribution is extremely complex for high dimensional data.
- How to build a complex model to approximate the data distribution?
 - start small !

Generative Model Categorization

- Based on representation of probability distribution:
 - existing generative modeling techniques → two categories
 - *Likelihood-based models*:
 - directly learns the distribution's prob. density fn via max. likelihood → VAEs
 - limitations:
 - require strong restriction on the model architecture for tractable normalization
 - rely on surrogate objectives to approximate maximum likelihood training
 - *Implicit generative models*:
 - prob. dist. is represented by a model of its sampling process → GANs
 - limitations:
 - require adversarial training → unstable and lead to mode collapse

PDF Estimation

- Data: $\{x_1, x_2, \dots, x_n\}$
- Underlying data distribution: $P(x)$
- $f_\theta(x) \in \mathbb{R}$: a real valued function parameterized by θ

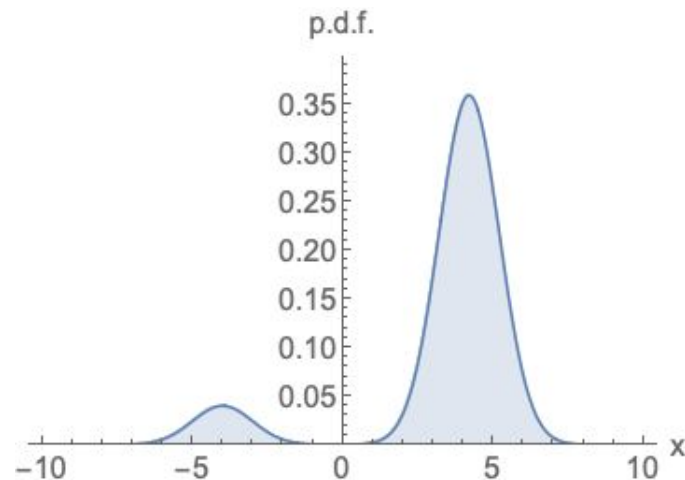
Probability Density Function (PDF): $P_\theta(x) = \frac{e^{f_\theta(x)}}{Z_\theta}$

Unnormalized Prob. model

Neural Network

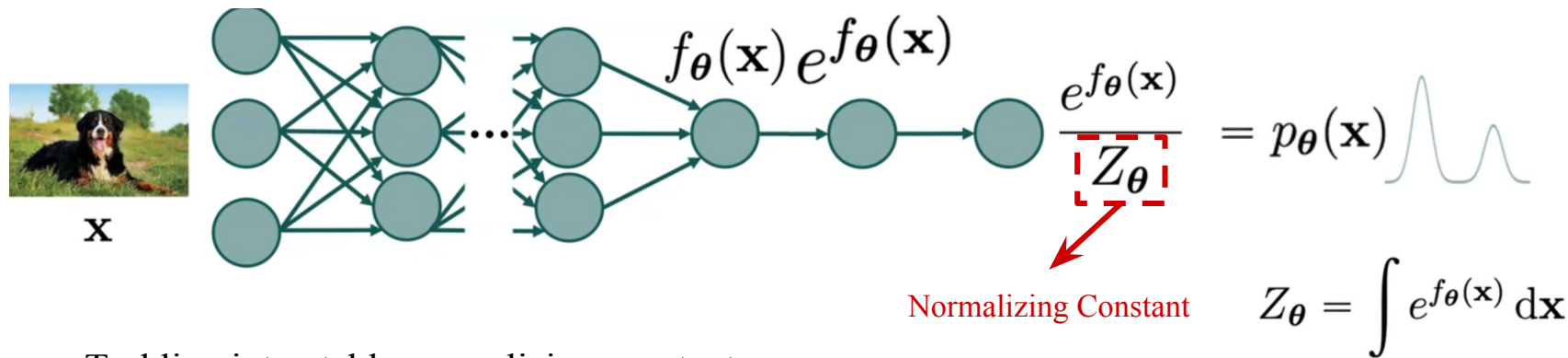
Normalizing constant

- Z is dependent on θ
- PDF should integrate to 1
- Parameterized PDF:
 - no matter how you change the model family and parameters, it has to be **normalized**.



PDF Estimation with DL Models

- DL model can approximate function in any range $(-\infty \text{ to } +\infty)$
- Probability density function is positive $\rightarrow$ require some modifications

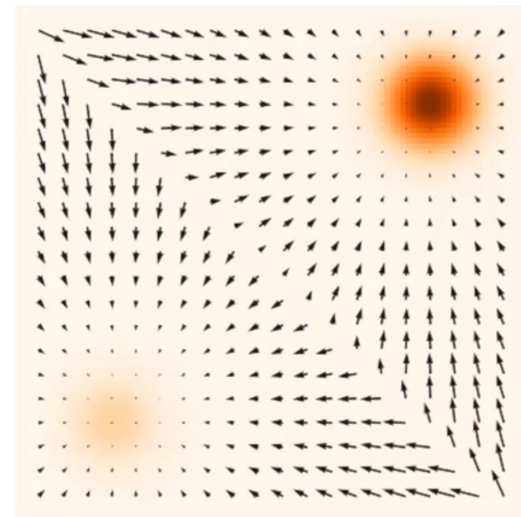


- Tackling intractable normalizing constant:
 - *Approximate normalizing constant:*
 - VAEs: leads to inaccurate probability evaluation
 - *Using restricted neural network models:*
 - highly restricted model family $\rightarrow$ can't capture complex data distribution
 - *Generative adversarial networks:*
 - can not evaluate probabilities
- Alternative method $\rightarrow$ score-based models

Intractable for most DL models 😞

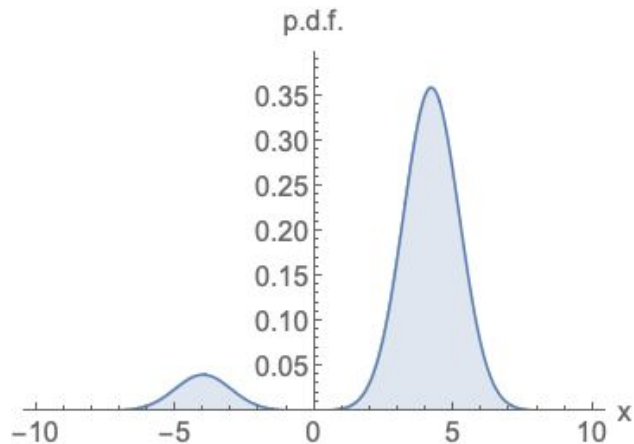
Score-Based Models

- **Proposed method: Score matching**
 - *model gradient of the log probability density function* $\rightarrow$ score function
 - do not require tractable normalizing constant
 - learning by score matching
 - Score function of a distribution $P(\mathbf{x}) = \nabla_{\mathbf{x}} \log p(\mathbf{x})$
- Score-based model: $\mathcal{S}_{\theta}(\mathbf{x}) \approx \nabla_{\mathbf{x}} \log p(\mathbf{x})$
 - gradient w.r.t $\mathbf{x}$, not θ
 - parameterized without normalizing constant
- **Score vs. Density function:**
 - density: two Gaussian distributions – color coded
 - score: vector field of gradients
 - points in the direction where density function has fastest growth
 - with derivative and integration: move from one to other

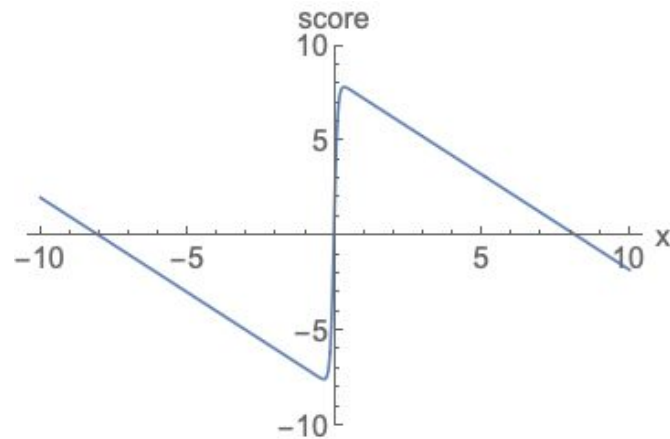


Score vs. density function

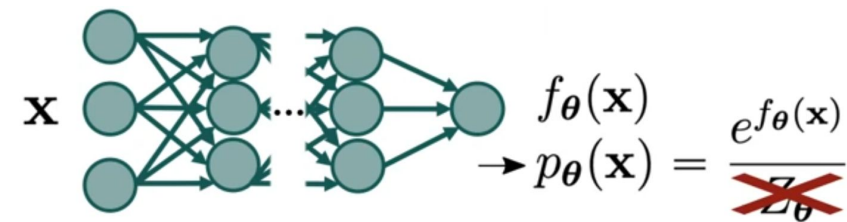
Score-Based Models



Probability Density Function



Score Function

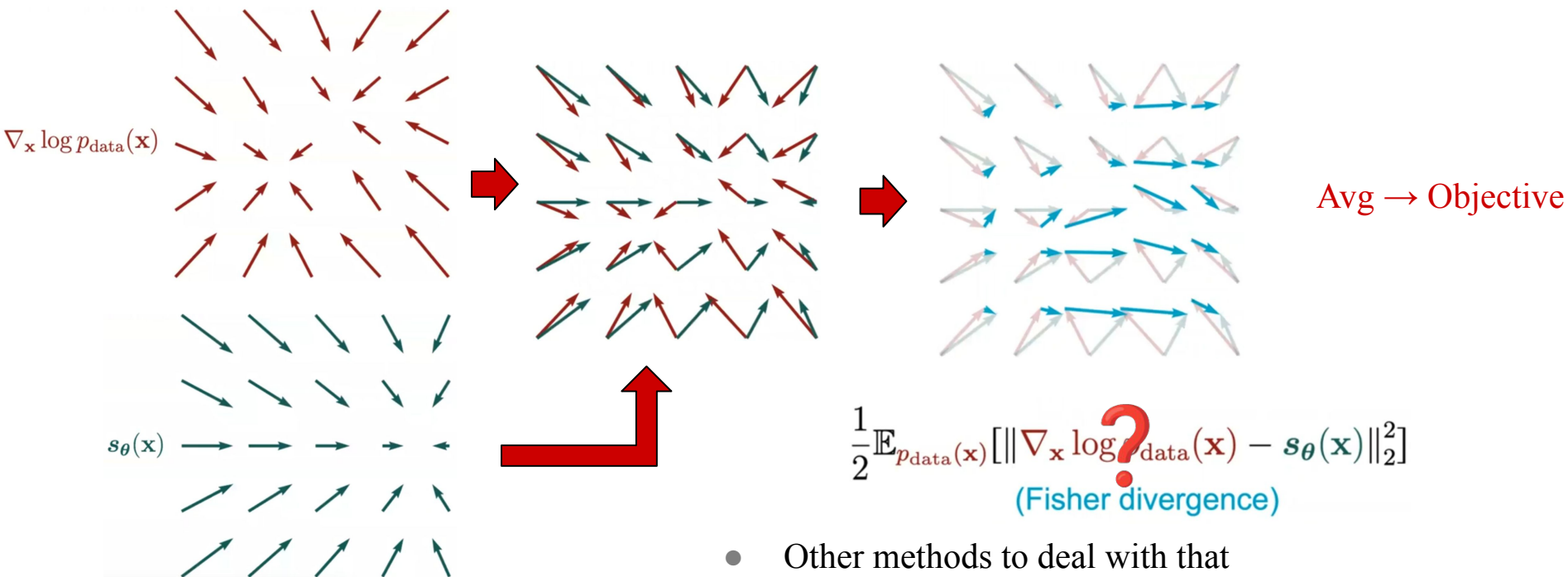


$$\begin{aligned}
 \nabla_{\mathbf{x}} \log p_{\theta}(\mathbf{x}) &= \nabla_{\mathbf{x}} f_{\theta}(\mathbf{x}) - \nabla_{\mathbf{x}} \log Z_{\theta} \\
 &= \nabla_{\mathbf{x}} f_{\theta}(\mathbf{x}) - \underset{\downarrow 0}{\boxed{\nabla_{\mathbf{x}} \log Z_{\theta}}}
 \end{aligned}$$

Training Score-Based Models

- **Given:** $\{x_1, x_2, \dots, x_n\} \stackrel{i.i.d.}{\sim} P_{data}(x)$
- **Goal:** $\nabla_{\mathbf{x}} \log p_{data}(\mathbf{x})$

- **Score Model:** $S_{\theta}(x) : R^d \rightarrow R^d \approx \nabla_{\mathbf{x}} \log p_{data}(\mathbf{x})$
- **Objective:** compare two vector fields. How?



- Other methods to deal with that
 - Sliced score matching (Song et al. 2019)

Score Matching

- Family of method that minimizes Fisher Divergence (FD):
 - without knowledge of the ground truth data score
 - optimized with stochastic gradient descent
 - do not require adversarial optimization
- Modeling flexibility of FD:
 - $S_{\theta}(x)$ doesn't have to be an actual score fn. of any normalized distribution
 - compares L2 distance b/w ground-truth data score and estimated score
 - no restrictions on the form or complexity of $S_{\theta}(x)$
 - it should only be a vector valued function
- Train score-based model with score matching, but how to generate new samples?

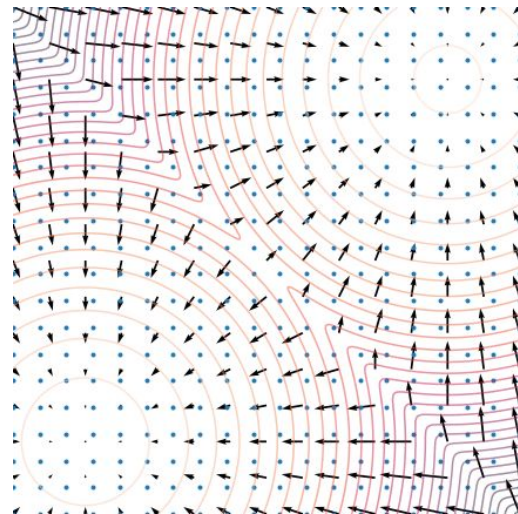
Sampling – Langevin Dynamics

- An iterative process to draw samples from a trained score-based model $S_{\theta}(x)$
- It provides an MCMC procedure to samples from a distribution $P(x)$ using only score function.
- It initializes the chain from an arbitrary prior distribution $x_0 \sim \pi(x)$
 - then iteratively samples

$$X_{i+1} \leftarrow X_i + \epsilon \nabla_{\mathbf{x}} \log p(\mathbf{x}) + \sqrt{2\epsilon} z_i, \quad i = 0, 1, \dots, K$$

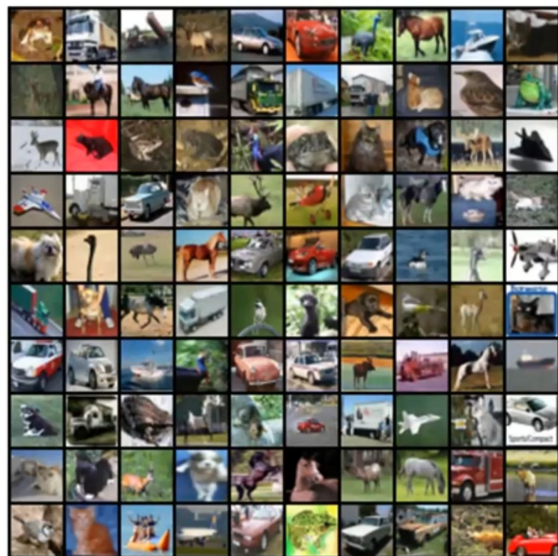
Noise $\sim N(0,1)$

- When $\epsilon \rightarrow 0$ and $K \rightarrow \infty$:
 - X_K obtained using above equation converges to a samples from $P(x)$
- In practice:
 - error is negligible when ϵ is sufficiently small and K is sufficiently large

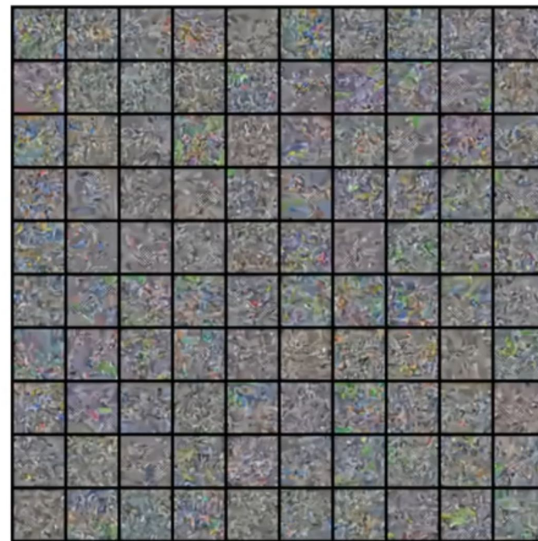


Naive Score-Based Generative Models

CIFAR-10 data



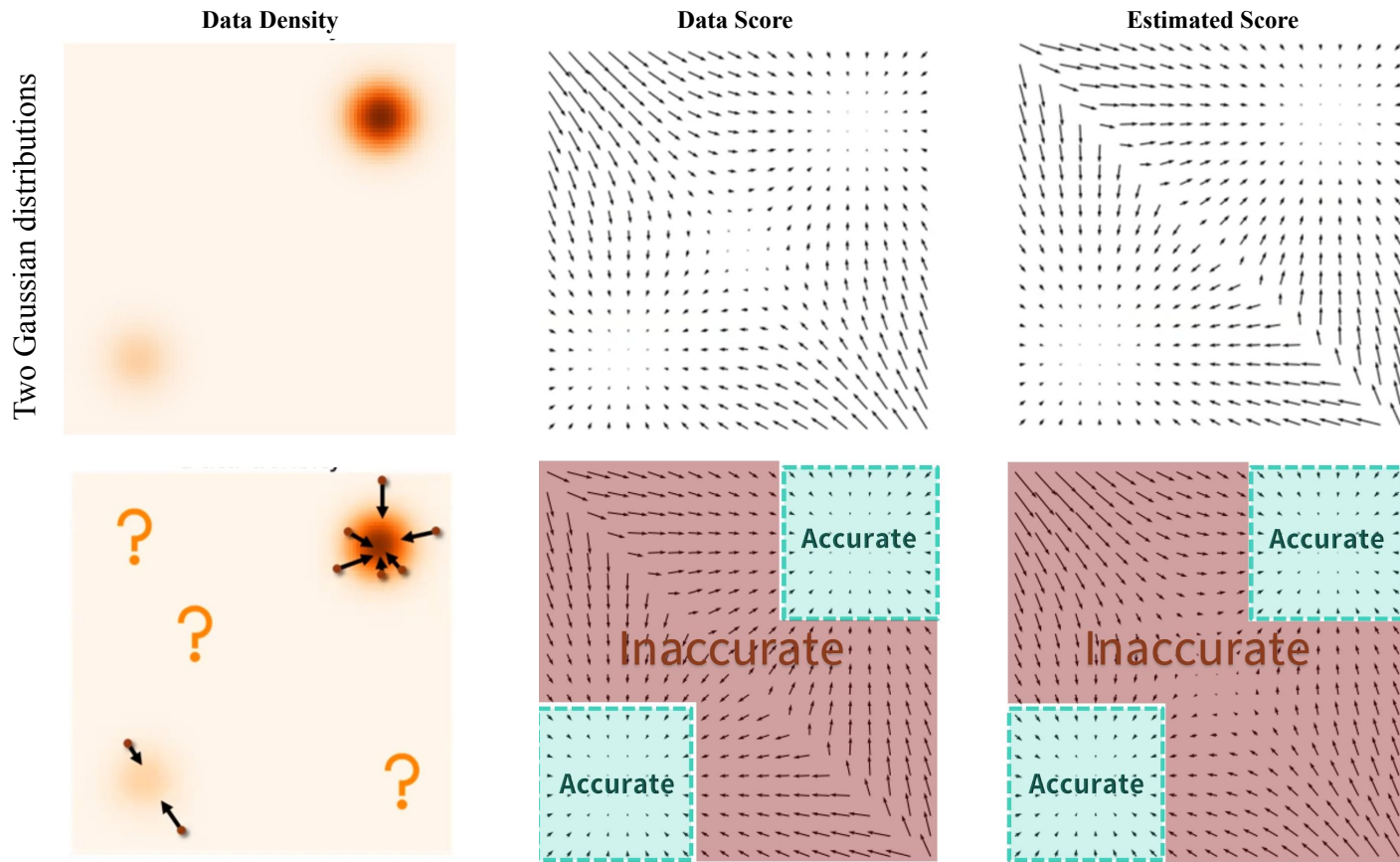
Model samples



- Score matching + Langevin dynamics $\Rightarrow$ doesn't work as expected !!
- **Reason:**
 - poor score estimation at low data density regions

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

Naive Score-Based Generative Models



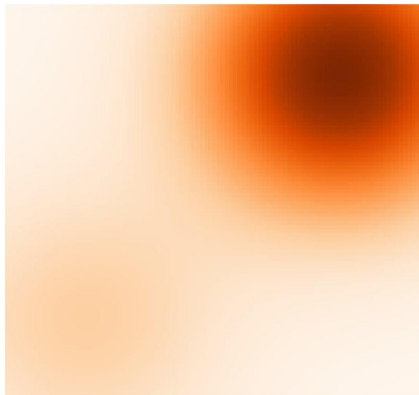
- Absence of sample in low data density regions makes the estimate inaccurate

Improving Score Estimation

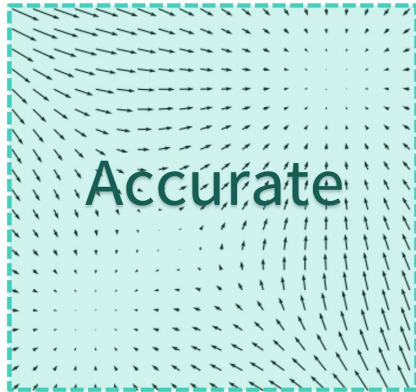
- Perturb data points with noise and train score-based models on the noisy data points
- When noise magnitude is sufficiently large $\rightarrow$ it can *populate low data density regions*
 - improve the accuracy of estimated score in that region
 - eg. mixture of Gaussians perturbed by additional Gaussian noise



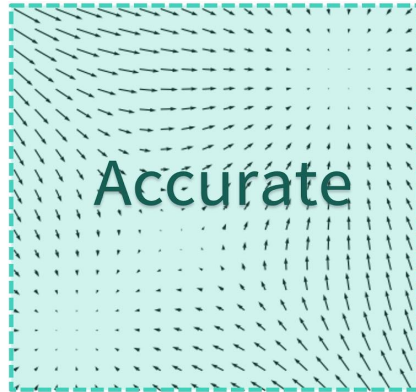
Perturbed Density



Perturbed Score



Estimated Score

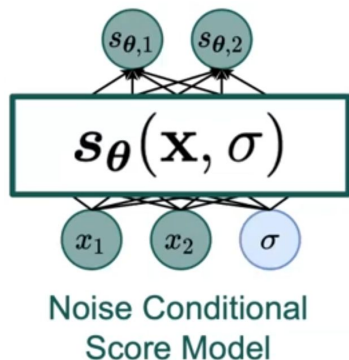


- Sequential application of noise
- How to choose an appropriate noise scale?



Choosing Noise Scale

- **Large noise:** cover more low density regions for better score estimation
 - but, it over corrupts the data $\rightarrow$ changing significantly from the original distribution
- **Low noise:** causes less corruption of the original data distribution
 - but, doesn't not cover the low density regions well !!
- Best of both world $\Rightarrow$ *Multiple scales of noise perturbations simultaneously*
 - perturbing with Gaussian noise with L increasing standard deviations: $\sigma_1 < \sigma_2 < \dots < \sigma_L$
 - *First:* perturb $P(x)$ with each of the Gaussian noise $N(0, \sigma_i^2 I)$, $i = 1, 2, \dots, L$
 - *Next:* estimate the score function of the each noisy-perturbed x by training
 - noise conditional score-based model $S_\theta(x, i)$ with score-matching, simultaneously



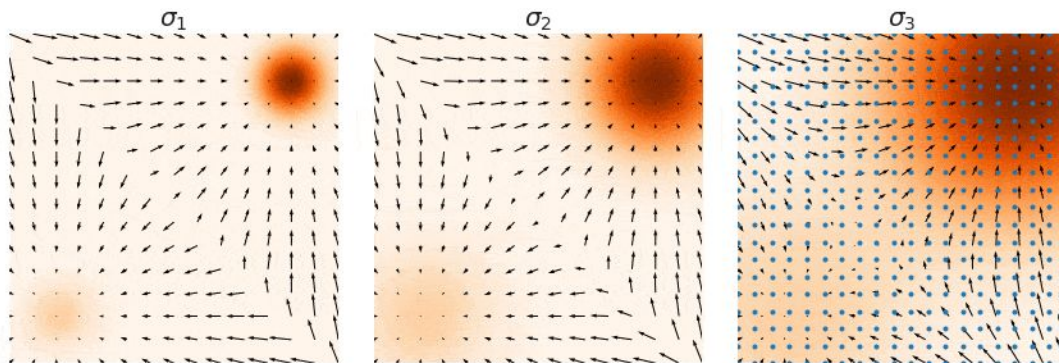
Positive weighting
function

$$\frac{1}{N} \sum_{i=1}^N \lambda(\sigma_i) \underbrace{\mathbb{E}_{p_{\sigma_i}(\mathbf{x})} [\|\nabla_{\mathbf{x}} \log p_{\sigma_i}(\mathbf{x}) - \mathbf{s}_\theta(\mathbf{x}, \sigma_i)\|_2^2]}_{\text{Score matching loss}}$$

Sampling – Annealed Langevin Dynamics

- **Annealed Langevin dynamics:**

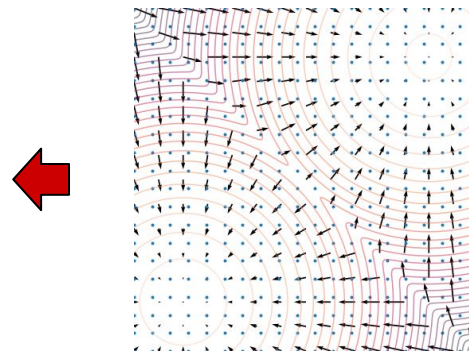
- combine a sequence of Langevin chains with gradually decreasing noise scales.
- Start with largest noise sample $\rightarrow$ slowly move towards smaller noise samples
- starting with largest noise model provides accurate estimate from the beginning
- slowly moving to smaller noise samples keeps this accuracy consistent



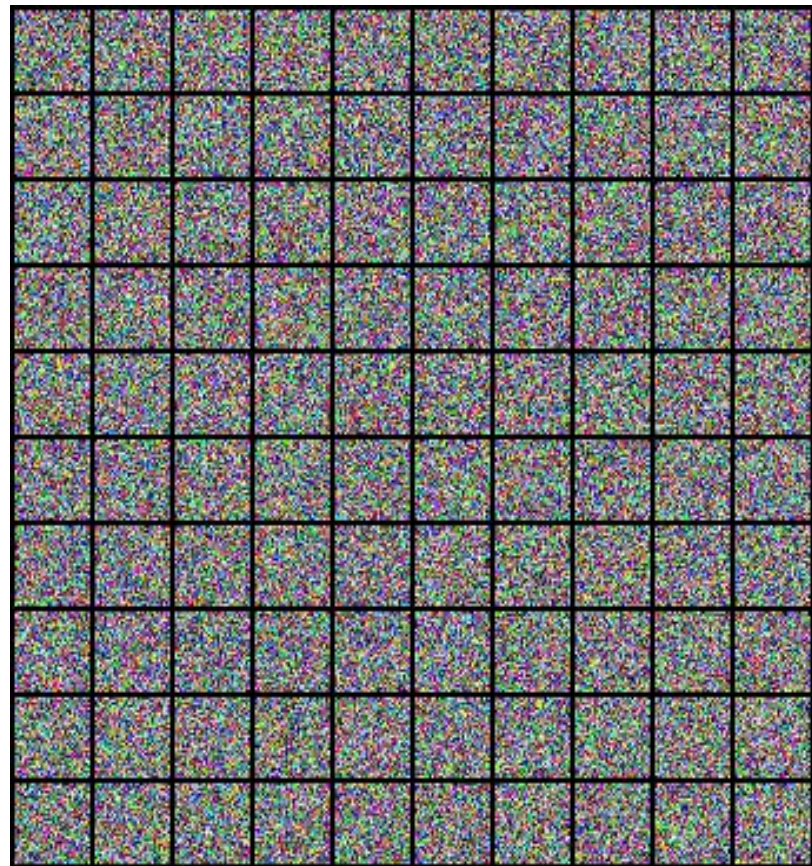
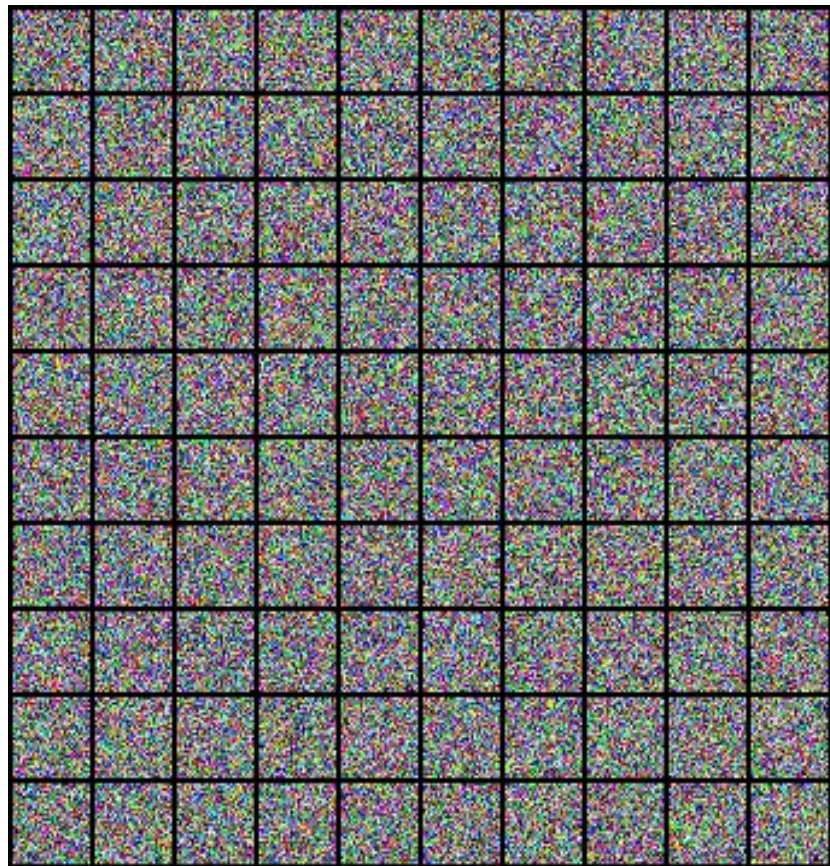
Algorithm 1 Annealed Langevin dynamics.

Require: $\{\sigma_i\}_{i=1}^L, \epsilon, T$.

```
1: Initialize  $\tilde{\mathbf{x}}_0$ 
2: for  $i \leftarrow 1$  to  $L$  do
3:    $\alpha_i \leftarrow \epsilon \cdot \sigma_i^2 / \sigma_L^2$   $\triangleright \alpha_i$  is the step size.
4:   for  $t \leftarrow 1$  to  $T$  do
5:     Draw  $\mathbf{z}_t \sim \mathcal{N}(0, I)$ 
6:      $\tilde{\mathbf{x}}_t \leftarrow \tilde{\mathbf{x}}_{t-1} + \frac{\alpha_i}{2} \mathbf{s}_\theta(\tilde{\mathbf{x}}_{t-1}, \sigma_i) + \sqrt{\alpha_i} \mathbf{z}_t$ 
7:   end for
8:    $\tilde{\mathbf{x}}_0 \leftarrow \tilde{\mathbf{x}}_T$ 
9: end for
return  $\tilde{\mathbf{x}}_T$ 
```



Samples Outputs



Diffusion Process - Intro

- **Goal:** learn structure of the probability density of a dataset
- Dye density represents probability density
 - each dye molecule is a sample
 - density of dye molecules represents probability density of data in that region
- **Observation:**
 - Diffusion process destroys structure in dataset



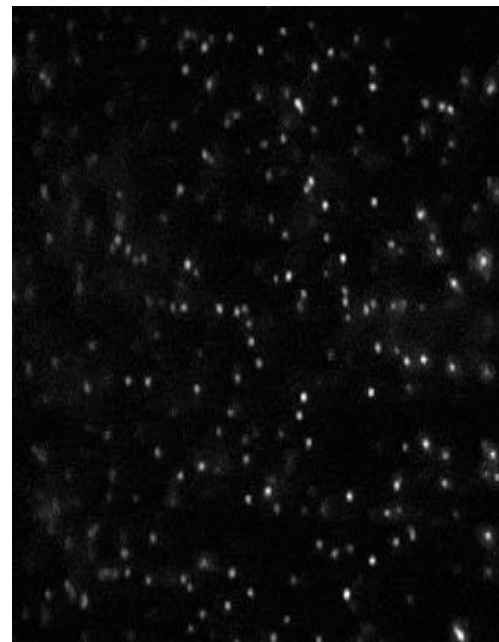
Diffusion Process - Intro

- **Core Idea:** Recover structure by reverse time
- Recover data distribution by
 - starting from uniform distribution and running backward dynamics
- Using Markov Chain, learn the Markov process for the reverse process exactly the way it was destroyed
- How exactly?
 - lets see diffusion more closely



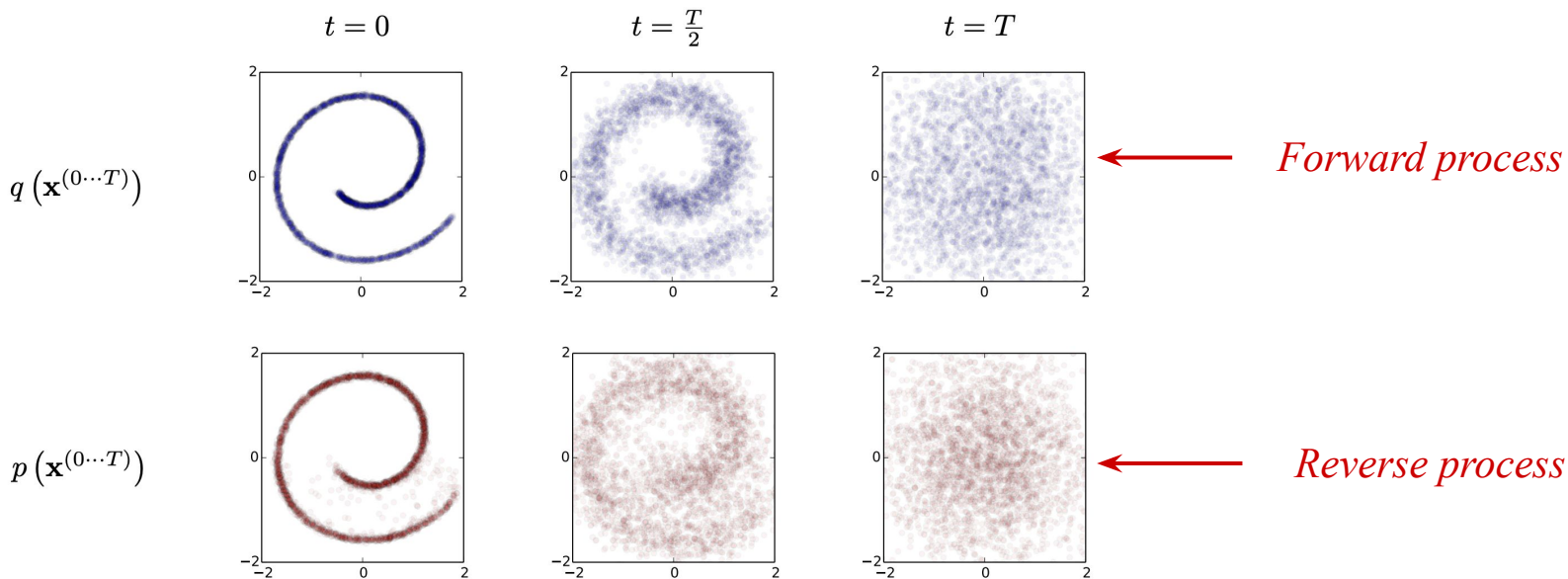
Microscopic View of Diffusion

- Brownian motion of dye molecules:
 - seems similar for forward and backward in time
- **Observation:**
 - Microscopic diffusion is time reversible
- Can we predict the next location of a molecule?
 - **Assume:** position updates are small Gaussian
 - Take a data point:
 - draw a next location from a small Gaussian centered around its current location
- For every small diffusion steps, the functional form of the reverse process is identical
- Reverse process also consists of drawing a Gaussian with same mean and small covariance close to the original data point



Algorithm

- Steps:
 - *Forward process*: Destroy all structure in data using diffusion process
 - symmetrically and slowly
 - *Reverse process*: Train Deep Learning model to reverse the diffusion process
 - Reverse diffusion process constitutes the generative model of data



Tractability Vs. Flexibility

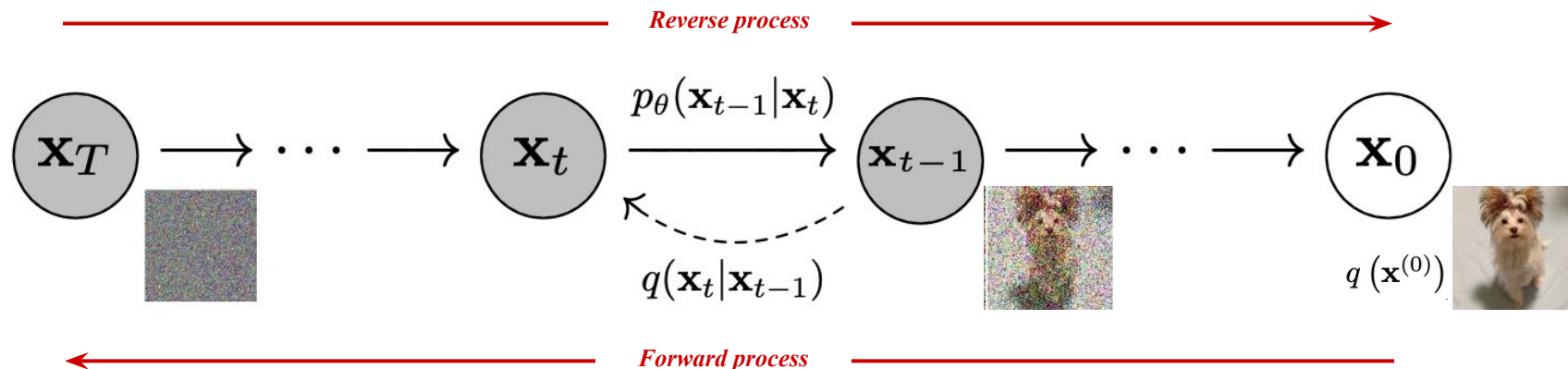
- *Tractable models* are easy to fit to data and can be analytically evaluated
 - but are unable to aptly describe structure in rich datasets
- *Flexible models* can arbitrarily fit any data
 - but are hard to evaluate
- Define a *flexible model* $\phi(x)$, yielding the flexible distribution $P(x) = \phi(x)/Z$
 - Z is the normalization constant $\rightarrow$ intractable to estimate
 - require specific process to estimate
- **Diffusion Probabilistic Models:**
 - Markov Chain gradually converts one distribution into another
 - extreme flexibility in model structure
 - exact sampling
- A generative Markov chain is learned to convert:
 - a *simple known distribution* $\rightarrow$ into a *target distribution (data distribution)*
 - the probabilistic model is defined as the endpoint of the Markov chain

Algorithm

- A forward process converts any complex data distribution into a tractable distribution



- Learn a finite-time reversal of this diffusion process → generative model



Mathematical Formulation

Tractable distribution $\pi(\mathbf{y}) = \int d\mathbf{y}' \overset{\text{Markov Diffusion Kernel}}{\boxed{T_{\pi}(\mathbf{y}|\mathbf{y}'; \beta)}} \overset{\text{Diffusion rate}}{\pi(\mathbf{y}')} \leftarrow$

$$q(\mathbf{x}^{(t)}|\mathbf{x}^{(t-1)}) = T_{\pi}(\mathbf{x}^{(t)}|\mathbf{x}^{(t-1)}; \beta_t)$$

T steps of diffusion $\Rightarrow q(\mathbf{x}^{(0 \dots T)}) = \boxed{q(\mathbf{x}^{(0)})} \prod_{t=1}^T q(\mathbf{x}^{(t)}|\mathbf{x}^{(t-1)})$ *Input data distribution*

$$p(\mathbf{x}^{(T)}) = \pi(\mathbf{x}^{(T)})$$

Tractable distribution

Finite time reversal $\Rightarrow p(\mathbf{x}^{(0 \dots T)}) = p(\mathbf{x}^{(T)}) \prod_{t=1}^T p(\mathbf{x}^{(t-1)}|\mathbf{x}^{(t)})$



Forward Process



Reverse Process

- Training amounts to maximizing the model log likelihood

Thank you!



INDIANA UNIVERSITY BLOOMINGTON