

COMP 690-204 Deep Generative Models

Lecture 4: Variational Autoencoder (VAE)

Jason Ren

Fall 2026, CS, UNC

**Acknowledgement: Many of the following slides are modified from Kaiming He (MIT)*



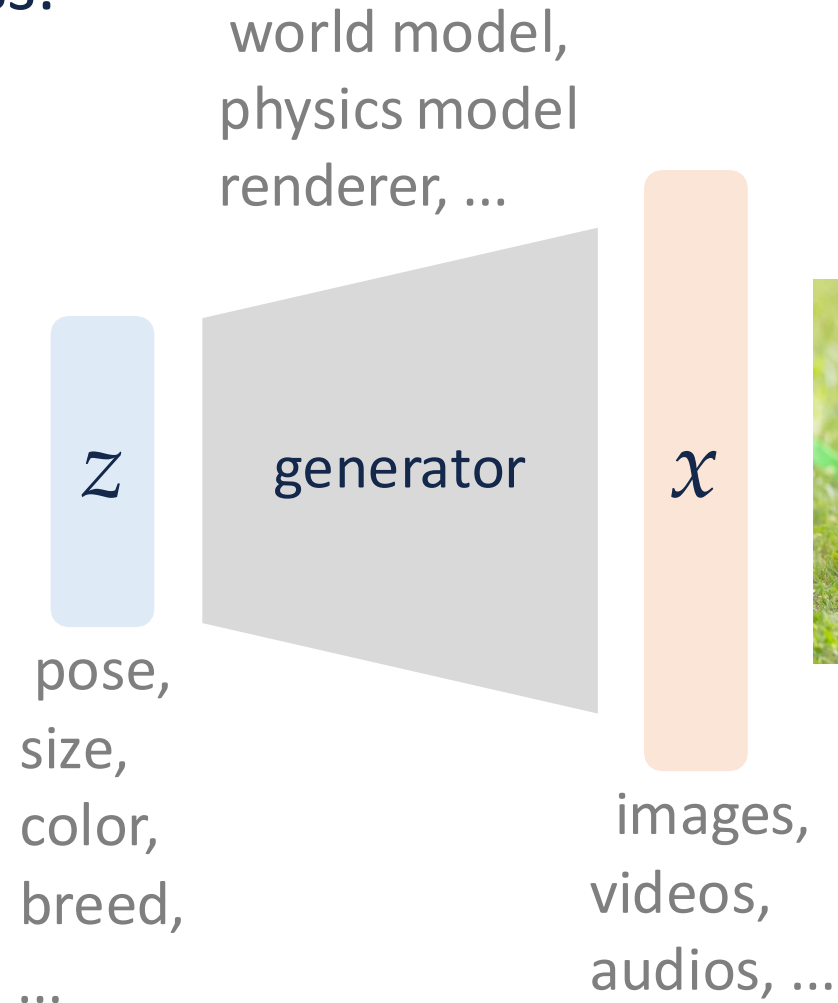
The University
of North Carolina
at Chapel Hill

Latent Variable Models



Assuming a data generation process:

- z - latent variables
- x - observed variables



Latent Variable Models

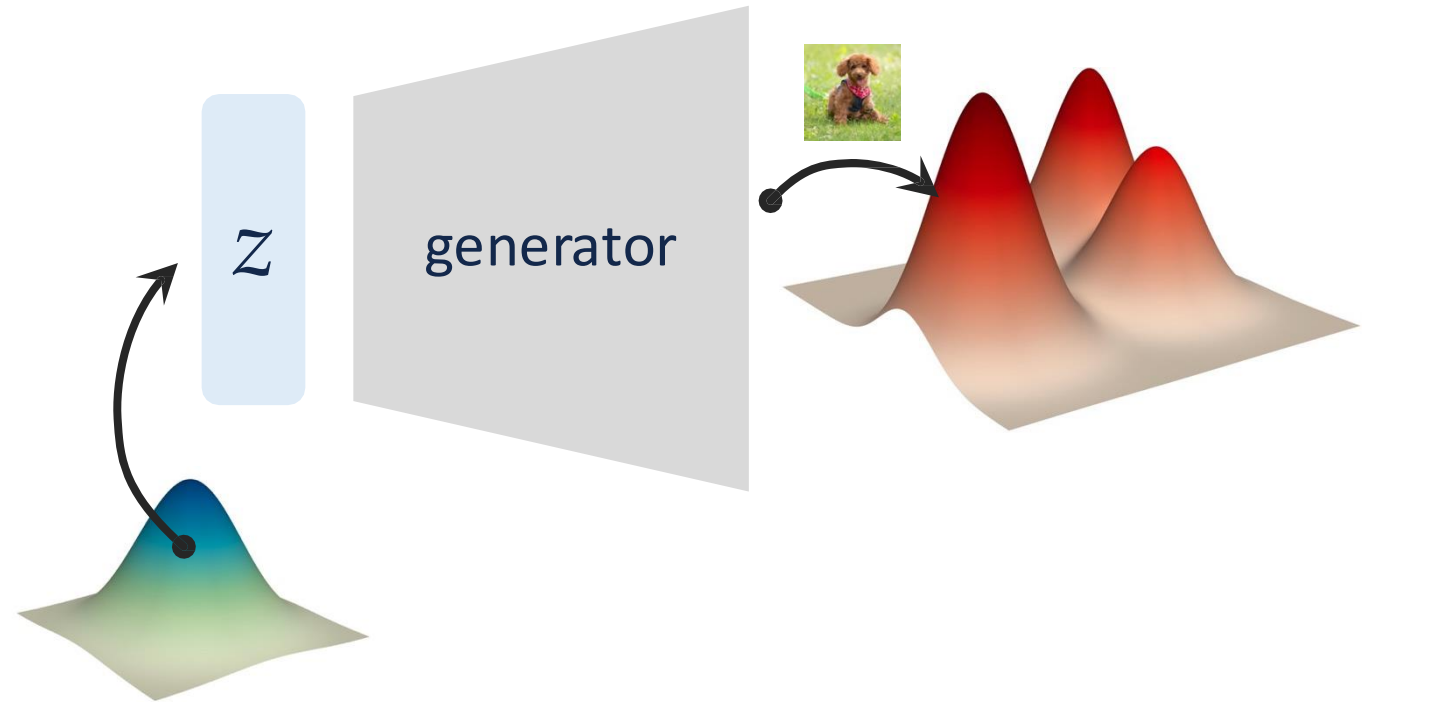


Assuming a data generation process:

- z - latent variables
- x - observed variables

sample from
prior distribution

$$z \sim p(z)$$

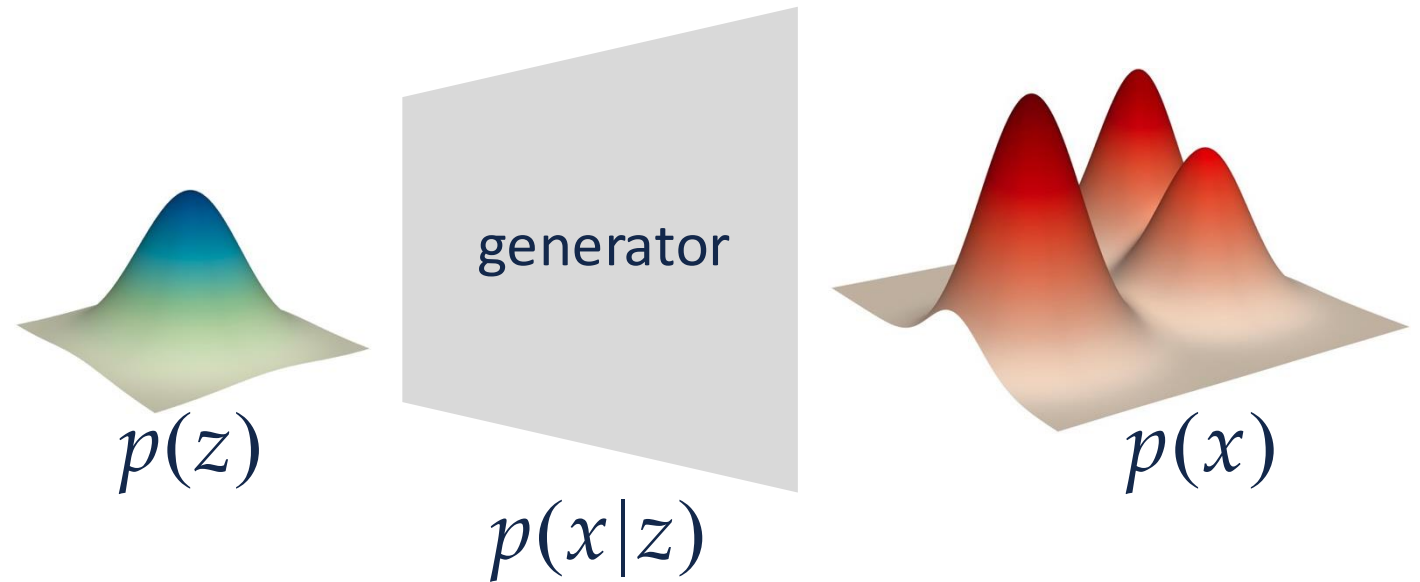


Latent Variable Models



Assuming a data generation process:

- z - latent variables
- x - observed variables

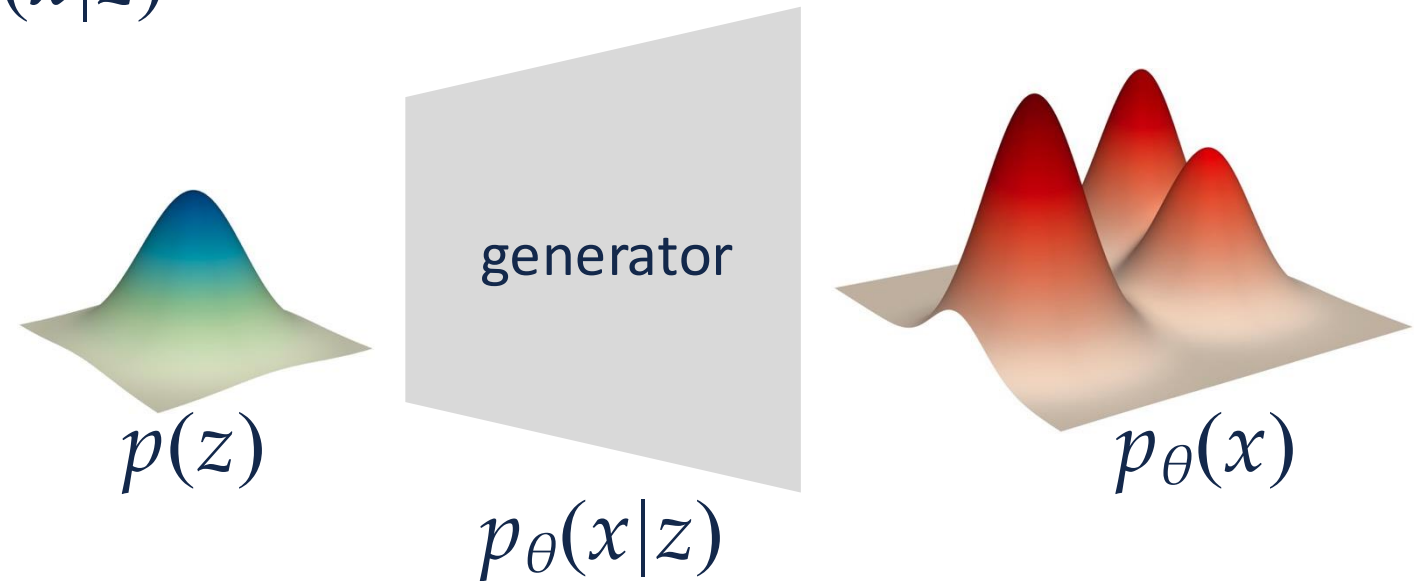


Latent Variable Models



Represent a distribution by a neural network

- θ - learnable parameters
- represent a function: $p_{\theta}(x|z)$



Measuring how good a distribution is ...



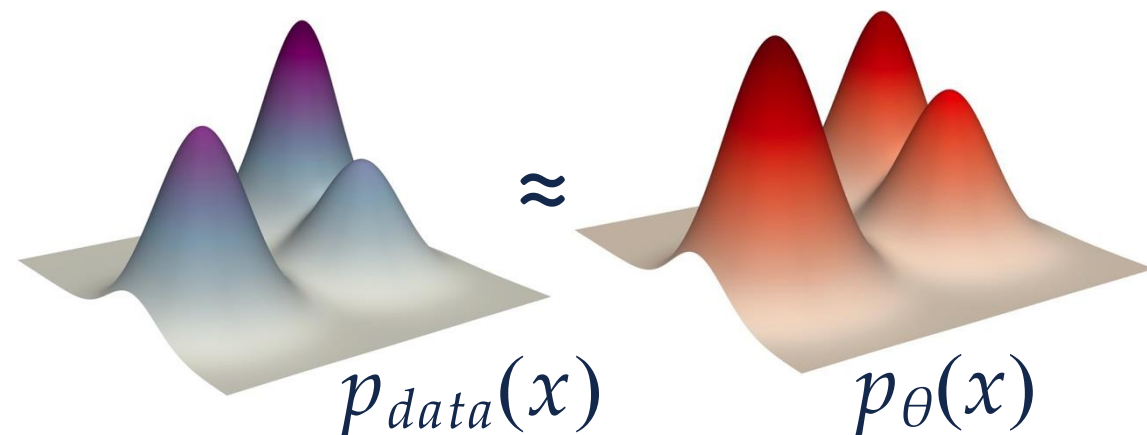
Minimize Kullback–Leibler (KL) divergence:

$$\min_{\theta} \mathcal{D}_{\text{KL}}(p_{\text{data}} \parallel p_{\theta})$$

⇒ Maximize likelihood:

$$\max_{\theta} \mathbb{E}_{x \sim p_{\text{data}}} \log p_{\theta}(x)$$

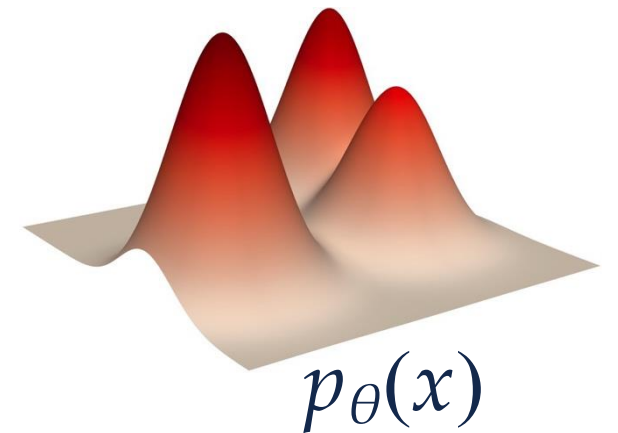
$$\begin{aligned} & \arg \min_{\theta} \mathcal{D}_{\text{KL}}(p_{\text{data}} \parallel p_{\theta}) && \text{tl; dr} \\ = & \arg \min_{\theta} \sum_x p_{\text{data}}(x) \log \frac{p_{\text{data}}(x)}{p_{\theta}(x)} \\ = & \arg \min_{\theta} \sum_x -p_{\text{data}}(x) \log p_{\theta}(x) + \text{const} \\ = & \arg \max_{\theta} \sum_x p_{\text{data}}(x) \log p_{\theta}(x) \\ = & \arg \max_{\theta} \mathbb{E}_{x \sim p_{\text{data}}} \log p_{\theta}(x) \end{aligned}$$



Latent Variable Models



We want to maximize $\mathbb{E}_{x \sim p_{data}} \log p_{\theta}(x)$

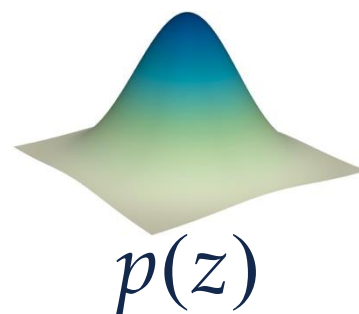


Latent Variable Models



We want to maximize $\mathbb{E}_{x \sim p_{data}} \log p_{\theta}(x)$
with $p_{\theta}(x)$ represented as:

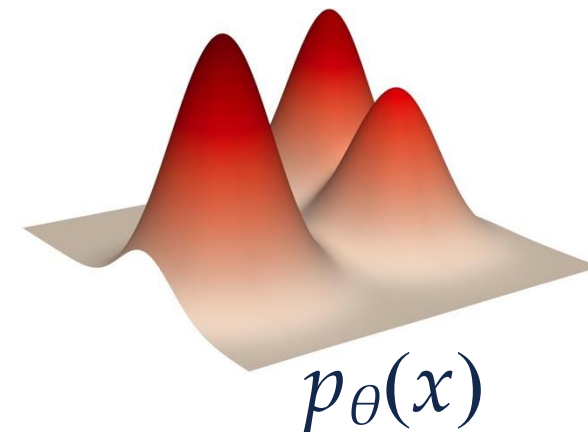
$$p_{\theta}(x) = \int_z p_{\theta}(x|z)p(z)dz$$



$p(z)$



$p_{\theta}(x|z)$



$p_{\theta}(x)$

Latent Variable Models

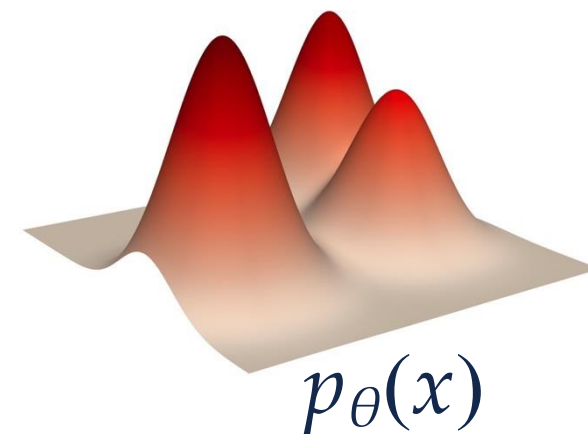
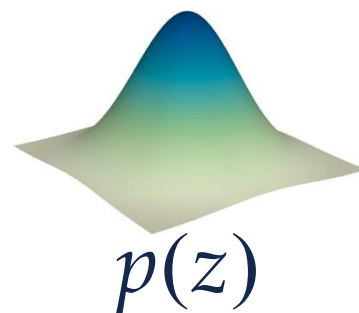


We want to maximize $\mathbb{E}_{x \sim p_{data}} \log p_{\theta}(x)$
with $p_{\theta}(x)$ represented as:

$$p_{\theta}(x) = \int_z \boxed{p_{\theta}(x|z)} \boxed{p(z)} dz$$

Two sets of unknowns:

- We need to optimize for θ
- We can't control **“true”** $p(z)$



Idea: introduce a “controllable” distribution $q(z)$

Latent Variable Models



$$\begin{aligned} & \log p_{\theta}(x) \\ &= \int_z q(z) \log p_{\theta}(x) dz \\ &= \int_z q(z) \log \left(\frac{p_{\theta}(x|z)p_{\theta}(z)}{p_{\theta}(z|x)} \right) dz \\ &= \int_z q(z) \log \left(\frac{p_{\theta}(x|z)p_{\theta}(z)}{p_{\theta}(z|x)} \frac{q(z)}{q(z)} \right) dz \\ &= \int_z q(z) \left(\log p_{\theta}(x|z) + \log \frac{p_{\theta}(z)}{q(z)} + \log \frac{q(z)}{p_{\theta}(z|x)} \right) dz \\ &= \mathbb{E}_{z \sim q(z)} \left[\log p_{\theta}(x|z) \right] - \mathcal{D}_{\text{KL}} \left(q(z) || p_{\theta}(z) \right) + \mathcal{D}_{\text{KL}} \left(q(z) || p_{\theta}(z|x) \right) \end{aligned}$$

Rewrite log likelihood by latent z

- for any distribution $q(z)$
- Bayes' rule

Latent Variable Models



$$\begin{aligned} & \log p_{\theta}(x) \\ = & \int_z q(z) \log p_{\theta}(x) dz \\ = & \int_z q(z) \log \left(\frac{p_{\theta}(x|z)p_{\theta}(z)}{p_{\theta}(z|x)} \right) dz \\ = & \int_z q(z) \log \left(\frac{p_{\theta}(x|z)p_{\theta}(z)}{p_{\theta}(z|x)} \frac{q(z)}{q(z)} \right) dz \\ = & \int_z q(z) \left(\log p_{\theta}(x|z) + \log \frac{p_{\theta}(z)}{q(z)} + \log \frac{q(z)}{p_{\theta}(z|x)} \right) dz \\ = & \mathbb{E}_{z \sim q(z)} \left[\log p_{\theta}(x|z) \right] - \mathcal{D}_{\text{KL}} \left(q(z) || p_{\theta}(z) \right) + \mathcal{D}_{\text{KL}} \left(q(z) || p_{\theta}(z|x) \right) \end{aligned}$$

Rewrite log likelihood by latent z

- for any distribution $q(z)$

- Bayes' rule

- just algebra

- just algebra

Latent Variable Models



$$\begin{aligned} & \log p_{\theta}(x) \\ = & \int_z q(z) \log p_{\theta}(x) dz \\ = & \int_z q(z) \log \left(\frac{p_{\theta}(x|z)p_{\theta}(z)}{p_{\theta}(z|x)} \right) dz \\ = & \int_z q(z) \log \left(\frac{p_{\theta}(x|z)p_{\theta}(z)}{p_{\theta}(z|x)} \frac{q(z)}{q(z)} \right) dz & \bullet \text{ just algebra} \\ = & \int_z q(z) \left(\log p_{\theta}(x|z) + \log \frac{p_{\theta}(z)}{q(z)} + \log \frac{q(z)}{p_{\theta}(z|x)} \right) dz & \bullet \text{ just algebra} \\ = & \mathbb{E}_{z \sim q(z)} \left[\log p_{\theta}(x|z) \right] - \mathcal{D}_{\text{KL}} \left(q(z) || p_{\theta}(z) \right) + \mathcal{D}_{\text{KL}} \left(q(z) || p_{\theta}(z|x) \right) \end{aligned}$$

Rewrite log likelihood by latent z

- for any distribution $q(z)$
- Bayes' rule

Latent Variable Models



$$\begin{aligned} & \log p_{\theta}(x) \\ = & \int_z q(z) \log p_{\theta}(x) dz \\ = & \int_z q(z) \log \left(\frac{p_{\theta}(x|z)p_{\theta}(z)}{p_{\theta}(z|x)} \right) dz \\ = & \int_z q(z) \log \left(\frac{p_{\theta}(x|z)p_{\theta}(z)}{p_{\theta}(z|x)} \frac{q(z)}{q(z)} \right) dz \\ = & \int_z q(z) \left(\log p_{\theta}(x|z) + \log \frac{p_{\theta}(z)}{q(z)} + \log \frac{q(z)}{p_{\theta}(z|x)} \right) dz \\ = & \mathbb{E}_{z \sim q(z)} \left[\log p_{\theta}(x|z) \right] - \mathcal{D}_{\text{KL}} \left(q(z) || p_{\theta}(z) \right) + \mathcal{D}_{\text{KL}} \left(q(z) || p_{\theta}(z|x) \right) \end{aligned}$$

Rewrite log likelihood by latent z

- for any distribution $q(z)$

- Bayes' rule

- just algebra

- just algebra

Latent Variable Models



$$\begin{aligned} & \log p_{\theta}(x) \\ = & \int_z q(z) \log p_{\theta}(x) dz \\ = & \int_z q(z) \log \left(\frac{p_{\theta}(x|z)p_{\theta}(z)}{p_{\theta}(z|x)} \right) dz \\ = & \int_z q(z) \log \left(\frac{p_{\theta}(x|z)p_{\theta}(z)}{p_{\theta}(z|x)} \frac{q(z)}{q(z)} \right) dz \\ = & \int_z q(z) \left(\log p_{\theta}(x|z) + \log \frac{p_{\theta}(z)}{q(z)} + \log \frac{q(z)}{p_{\theta}(z|x)} \right) dz \\ = & \mathbb{E}_{z \sim q(z)} \left[\log p_{\theta}(x|z) \right] - \mathcal{D}_{\text{KL}} \left(q(z) || p_{\theta}(z) \right) + \mathcal{D}_{\text{KL}} \left(q(z) || p_{\theta}(z|x) \right) \end{aligned}$$

Rewrite log likelihood by latent z

- for any distribution $q(z)$

- Bayes' rule

- just algebra

- just algebra

Latent Variable Models



$$\begin{aligned} & \log p_{\theta}(x) \\ = & \int_z q(z) \log p_{\theta}(x) dz \\ = & \int_z q(z) \log \left(\frac{p_{\theta}(x|z)p_{\theta}(z)}{p_{\theta}(z|x)} \right) dz \\ = & \int_z q(z) \log \left(\frac{p_{\theta}(x|z)p_{\theta}(z)}{p_{\theta}(z|x)} \frac{q(z)}{q(z)} \right) dz \\ = & \int_z q(z) \left(\log p_{\theta}(x|z) + \log \frac{p_{\theta}(z)}{q(z)} + \log \frac{q(z)}{p_{\theta}(z|x)} \right) dz \\ = & \mathbb{E}_{z \sim q(z)} \left[\log p_{\theta}(x|z) \right] - \mathcal{D}_{\text{KL}} \left(q(z) || p_{\theta}(z) \right) + \mathcal{D}_{\text{KL}} \left(q(z) || p_{\theta}(z|x) \right) \end{aligned}$$

Rewrite log likelihood by latent z

- for any distribution $q(z)$

- Bayes' rule

- just algebra

- just algebra

Latent Variable Models



intractable

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

$$= \int_z q(z) \log p_{\theta}(x) dz$$

$$= \int_z q(z) \log \left(\frac{p_{\theta}(x|z)p_{\theta}(z)}{p_{\theta}(z|x)} \right) dz$$

$$= \int_z q(z) \log \left(\frac{p_{\theta}(x|z)p_{\theta}(z)}{p_{\theta}(z|x)} \frac{q(z)}{q(z)} \right) dz$$

$$= \int_z q(z) \left(\log p_{\theta}(x|z) + \log \frac{p_{\theta}(z)}{q(z)} + \log \frac{q(z)}{p_{\theta}(z|x)} \right) dz$$

$$= \mathbb{E}_{z \sim q(z)} \left[\log p_{\theta}(x|z) \right] - \mathcal{D}_{\text{KL}} \left(q(z) || p_{\theta}(z) \right) + \mathcal{D}_{\text{KL}} \left(q(z) || p_{\theta}(z|x) \right)$$

tractable

tractable

intractable

Rewrite log likelihood by latent z

- for any distribution $q(z)$
- Bayes' rule

Latent Variable Models



$$\begin{aligned} & \text{intractable} \quad \boxed{\log p_{\theta}(x)} - \boxed{\mathcal{D}_{\text{KL}}(q(z) || p_{\theta}(z|x))} \quad \text{intractable} \\ &= \underbrace{\mathbb{E}_{z \sim q(z)} [\log p_{\theta}(x|z)]}_{\text{tractable}} - \underbrace{\mathcal{D}_{\text{KL}}(q(z) || p_{\theta}(z))}_{\text{tractable}} \end{aligned}$$

Latent Variable Models



intractable $\log p_\theta(x) - \mathcal{D}_{\text{KL}}(q(z) || p_\theta(z|x))$ intractable

$$= \underbrace{\mathbb{E}_{z \sim q(z)} [\log p_\theta(x|z)]}_{\text{tractable}} - \underbrace{\mathcal{D}_{\text{KL}}(q(z) || p_\theta(z))}_{\text{tractable}}$$

- This is called Evidence Lower Bound (ELBO)
- Lower bound of $\log p_\theta(x)$
- This equation holds for any distribution $q(z)$

Latent Variable Models



$$\underbrace{\mathbb{E}_{z \sim \cancel{q(z)}}^{\cancel{q_\varphi(z|x)}} \left[\log p_\theta(x|z) \right] - \mathcal{D}_{\text{KL}} \left(\cancel{q(z)} \parallel p_\theta(z) \right)}_{\text{tractable}}$$

- This is called Evidence Lower Bound (ELBO)
- Lower bound of $\log p_\theta(x)$
- This equation holds for any distribution $q(z)$
- Parameterize $q(z)$ by $q_\varphi(z|x)$

Latent Variable Models



$$\underbrace{\mathbb{E}_{z \sim \cancel{q(z)}}^{\cancel{q_\phi(z|x)}} \left[\log p_\theta(x|z) \right] - \mathcal{D}_{\text{KL}} \left(\cancel{q(z)} \parallel \cancel{p_\theta(z)} \right)}_{\text{tractable}}^{\cancel{q_\phi(z|x)} \quad p(z)}$$

- This is called Evidence Lower Bound (ELBO)
- Lower bound of $\log p_\theta(x)$
- This equation holds for any distribution $q(z)$
- Parameterize $q(z)$ by $q_\phi(z|x)$
- let $p_\theta(z)$ be a simple known prior $p(z)$

Variational Autoencoder



Maximize ELBO $\Rightarrow$ minimize an objective:

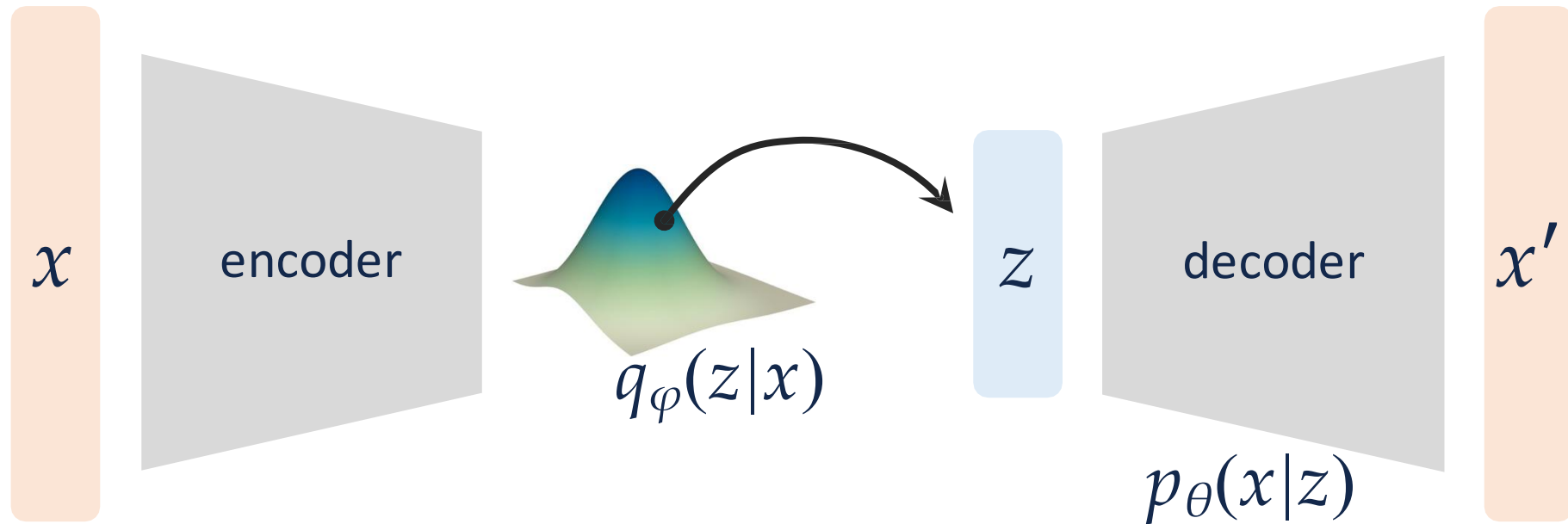
$$\mathcal{L}_{\theta, \phi}(x) = -\mathbb{E}_{z \sim q_{\phi}(z|x)} \left[\log p_{\theta}(x|z) \right] + \mathcal{D}_{\text{KL}} \left(q_{\phi}(z|x) || p(z) \right)$$

Variational Autoencoder



Maximize ELBO $\Rightarrow$ minimize an objective:

$$\mathcal{L}_{\theta, \phi}(x) = -\mathbb{E}_{z \sim q_{\phi}(z|x)} \left[\log p_{\theta}(x|z) \right] + \mathcal{D}_{\text{KL}} \left(q_{\phi}(z|x) || p(z) \right)$$

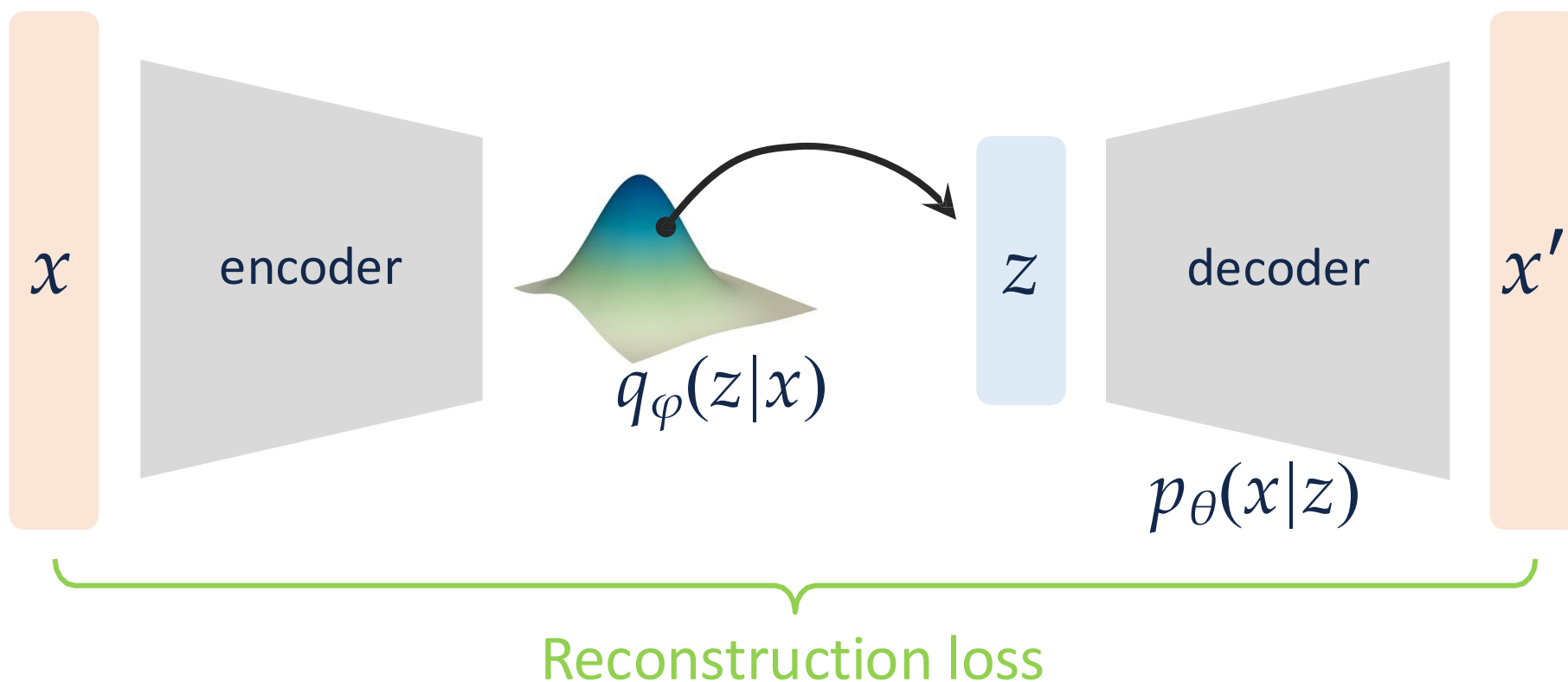


Variational Autoencoder



Maximize ELBO $\Rightarrow$ minimize an objective:

$$\mathcal{L}_{\theta, \phi}(x) = \boxed{-\mathbb{E}_{z \sim q_{\phi}(z|x)} \left[\log p_{\theta}(x|z) \right]} + \mathcal{D}_{\text{KL}} \left(q_{\phi}(z|x) || p(z) \right)$$

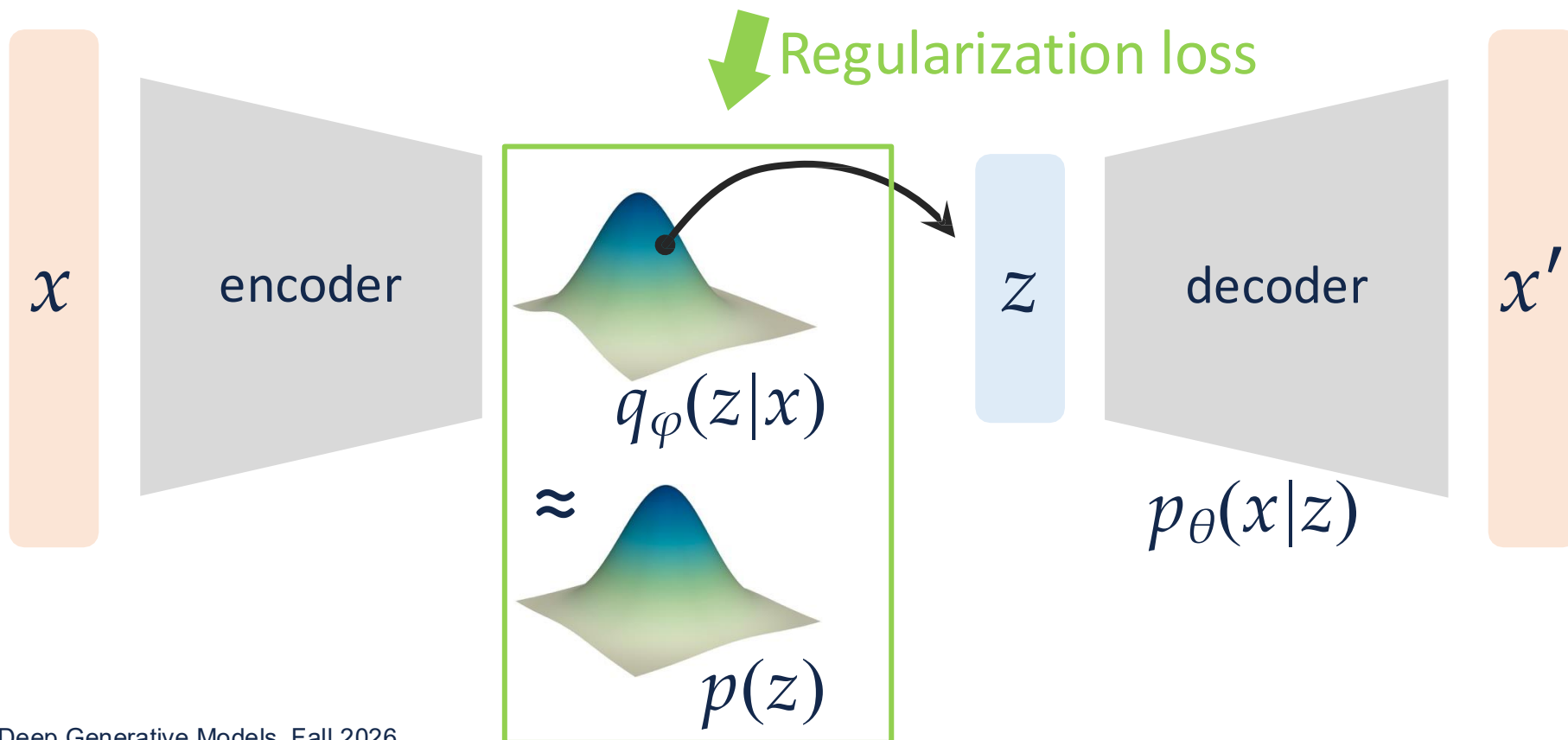


Variational Autoencoder



Maximize ELBO $\Rightarrow$ minimize an objective:

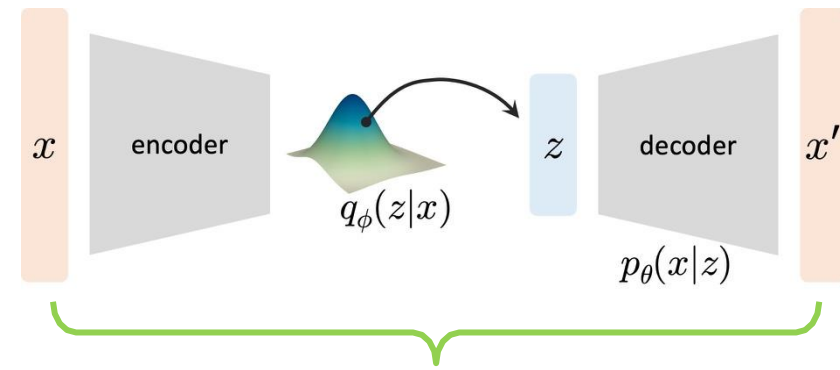
$$\mathcal{L}_{\theta, \phi}(x) = -\mathbb{E}_{z \sim q_{\phi}(z|x)} \left[\log p_{\theta}(x|z) \right] + \mathcal{D}_{\text{KL}} \left(q_{\phi}(z|x) || p(z) \right)$$



Variational Autoencoder

Reconstruction loss

$$\mathcal{L}_{\theta, \phi}(x) = \boxed{-\mathbb{E}_{z \sim q_{\phi}(z|x)} \left[\log p_{\theta}(x|z) \right]} + \mathcal{D}_{\text{KL}}(q_{\phi}(z|x) || p(z))$$

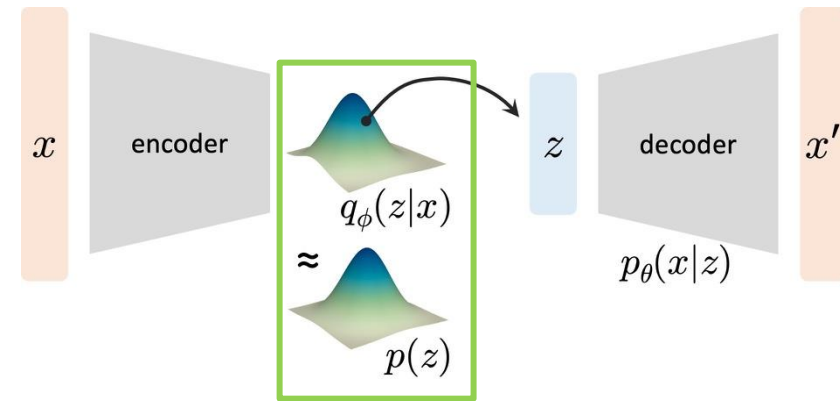


Example: L2 loss

- one-step Monte Carlo: $z \sim q_{\phi}(z|x)$
- map z by decoder net: $g_{\theta}(z) \rightarrow x'$ network estimates distribution's parameters
- model $p_{\theta}(x|z)$ by Gaussian: $p_{\theta}(x|z) = \mathcal{N}(x | x', \sigma_0^2)$ (assume fixed std)
- negative log likelihood: $\frac{1}{2\sigma_0^2} \|x - x'\|^2 + \text{const}$
- L2 loss $\Rightarrow$ a Gaussian neighborhood around data point x

Variational Autoencoder

Regularization loss



$$\mathcal{L}_{\theta, \phi}(x) = -\mathbb{E}_{z \sim q_\phi(z|x)} [\log p_\theta(x|z)] + \mathcal{D}_{\text{KL}}(q_\phi(z|x) || p(z))$$

Example: Gaussian prior

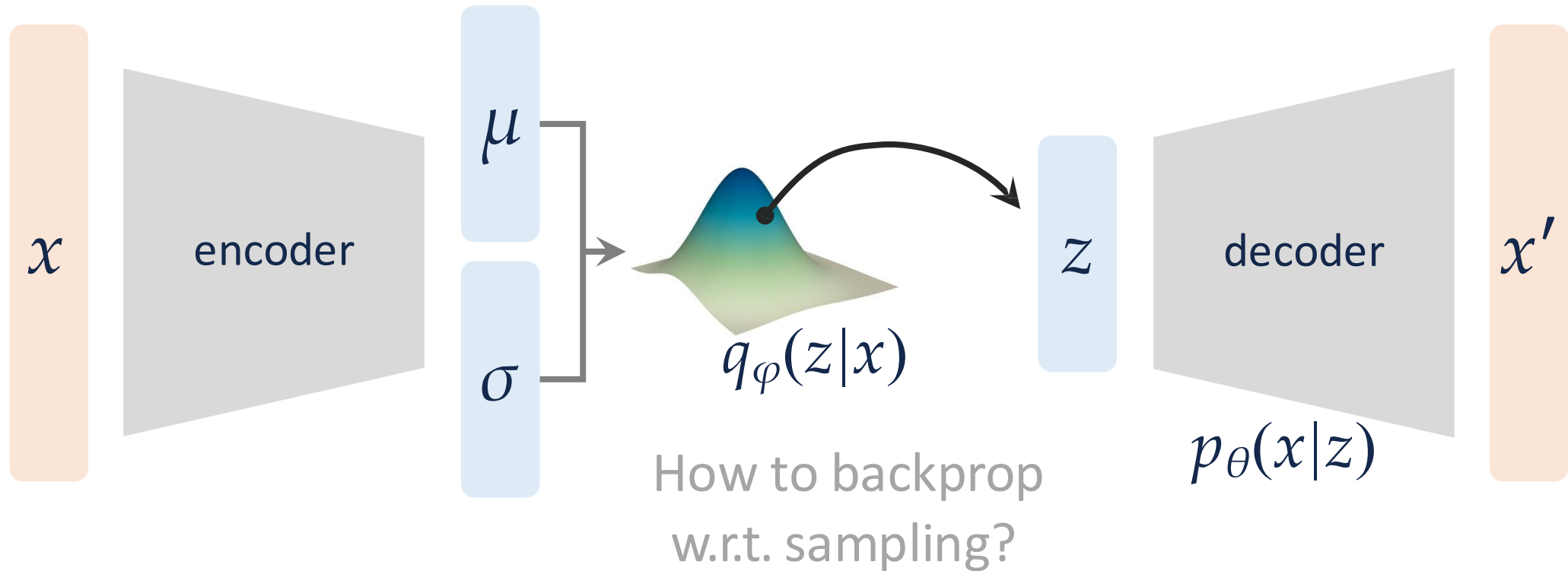
- let $p(z) = \mathcal{N}(z | 0, \mathbf{I})$
- model $q_\phi(z|x)$ by Gaussian: $\mathcal{N}(z | \mu, \sigma)$
- map x by encoder net: $f_\phi(x) \rightarrow \mu, \sigma$ again, network estimates distribution's parameters
- compute loss analytically: $\mathcal{D}_{\text{KL}}(\mathcal{N}(z | \mu, \sigma) || \mathcal{N}(z | 0, \mathbf{I}))$ (see pset 1)
- fixed covariance $\Rightarrow$ L2 loss on μ (see pset 1)

Variational Autoencoder



Maximize ELBO $\Rightarrow$ minimize an objective:

$$\mathcal{L}_{\theta, \phi}(x) = -\mathbb{E}_{z \sim q_{\phi}(z|x)} \left[\log p_{\theta}(x|z) \right] + \mathcal{D}_{\text{KL}} \left(q_{\phi}(z|x) || p(z) \right)$$

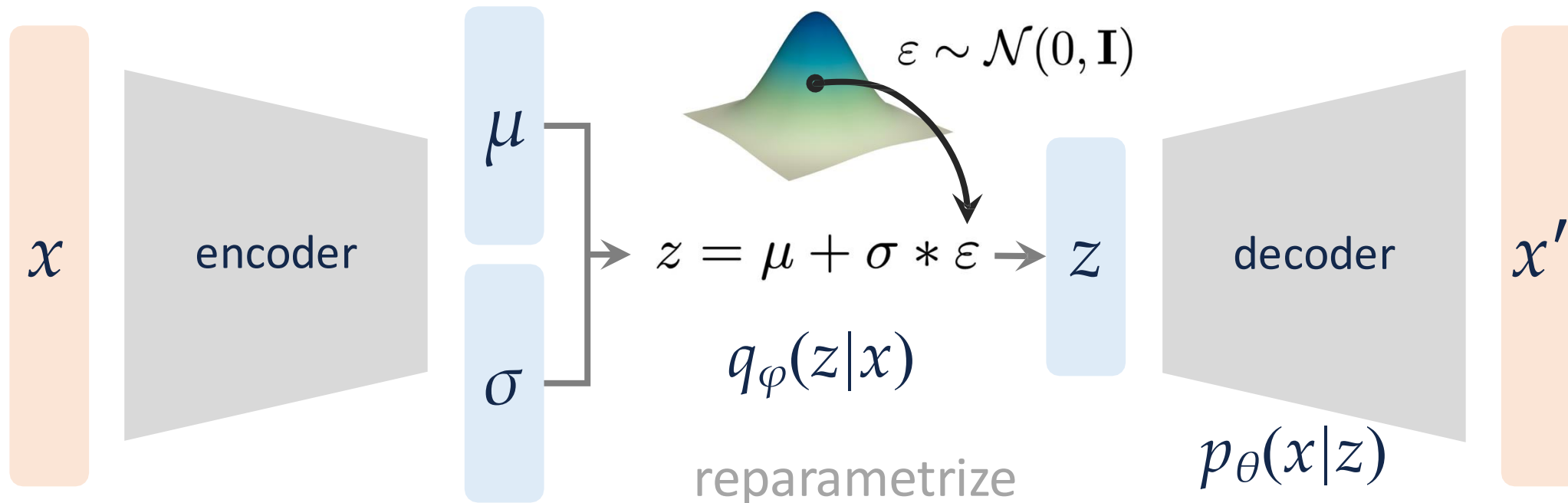


Variational Autoencoder



Maximize ELBO $\Rightarrow$ minimize an objective:

$$\mathcal{L}_{\theta, \phi}(x) = -\mathbb{E}_{z \sim q_{\phi}(z|x)} \left[\log p_{\theta}(x|z) \right] + \mathcal{D}_{\text{KL}} \left(q_{\phi}(z|x) || p(z) \right)$$



Variational Autoencoder



... so far, we have discussed an objective on one x :

$$\mathcal{L}_{\theta, \phi}(x) = \boxed{-\mathbb{E}_{z \sim q_{\phi}(z|x)} \left[\log p_{\theta}(x|z) \right]} + \boxed{\mathcal{D}_{\text{KL}} \left(q_{\phi}(z|x) || p(z) \right)}$$

Overall loss is expectation over data:

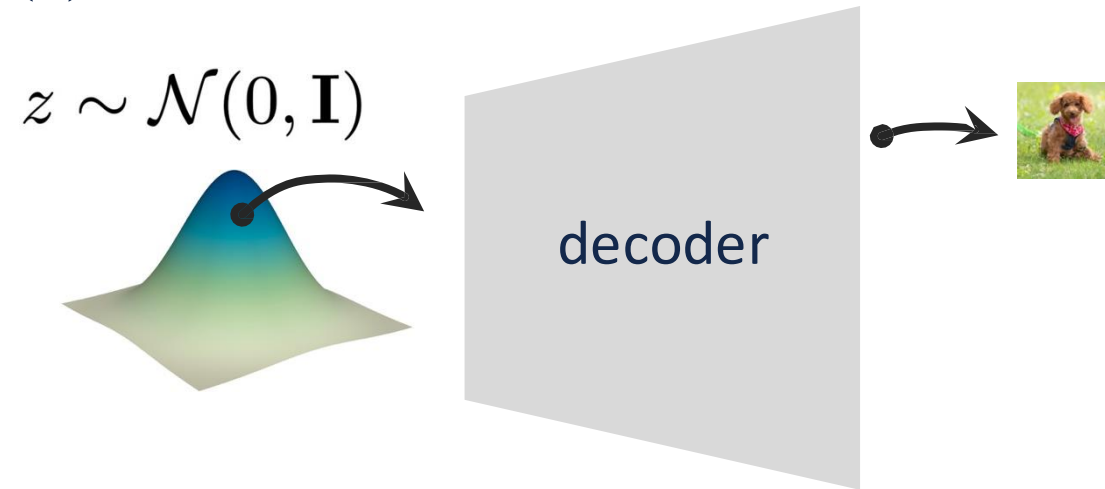
$$\mathcal{L}_{\theta, \phi} = \mathbb{E}_{x \sim p_{data}(x)} \left[-\mathbb{E}_{z \sim q_{\phi}(z|x)} \left[\log p_{\theta}(x|z) \right] + \mathcal{D}_{\text{KL}} \left(q_{\phi}(z|x) || p(z) \right) \right]$$

Variational Autoencoder



Inference (generation):

- sample z from: $\mathcal{N}(0, \mathbf{I})$
- map z by decoder net: $g_{\theta}(z)$

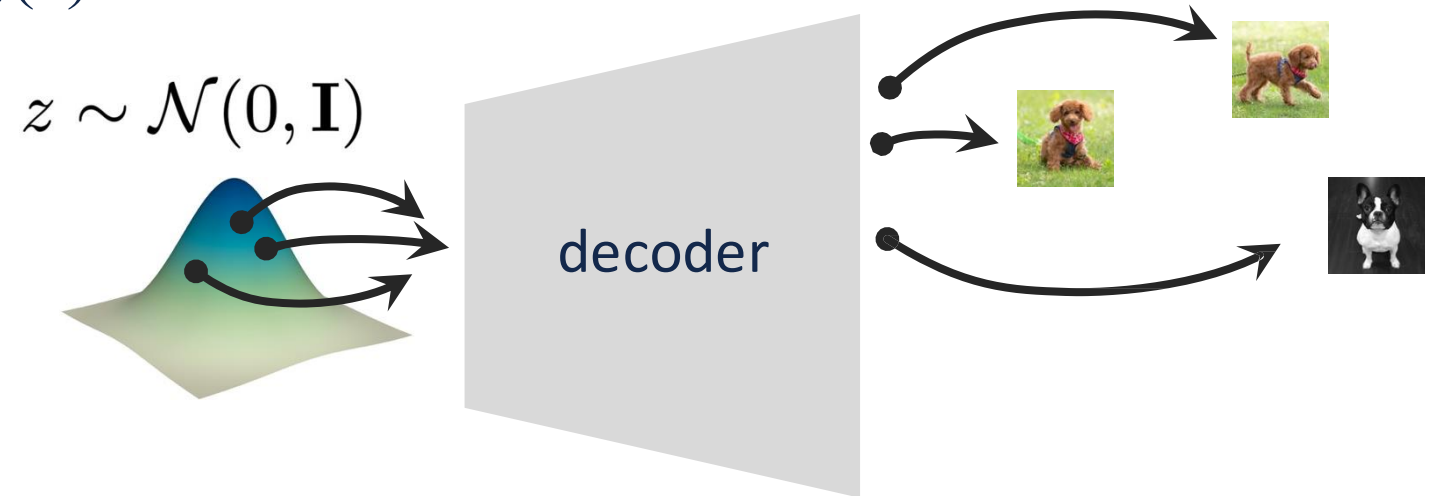


Variational Autoencoder



Inference (generation):

- sample z from: $\mathcal{N}(0, \mathbf{I})$
- map z by decoder net: $g_{\theta}(z)$

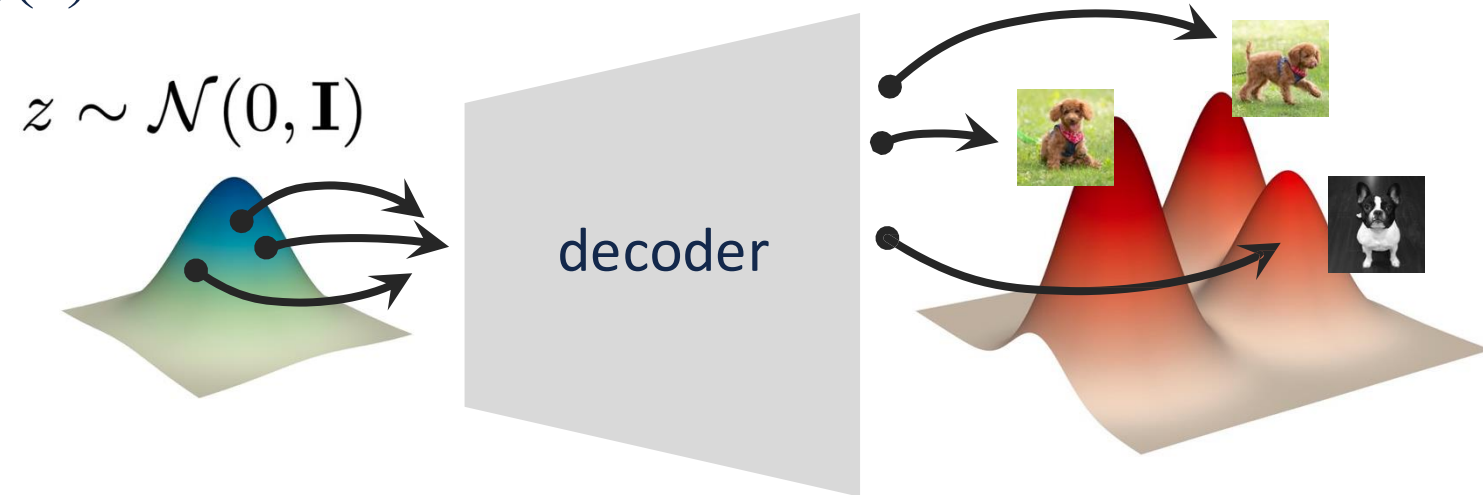


Variational Autoencoder



Inference (generation):

- sample z from: $\mathcal{N}(0, \mathbf{I})$
- map z by decoder net: $g_{\theta}(z)$

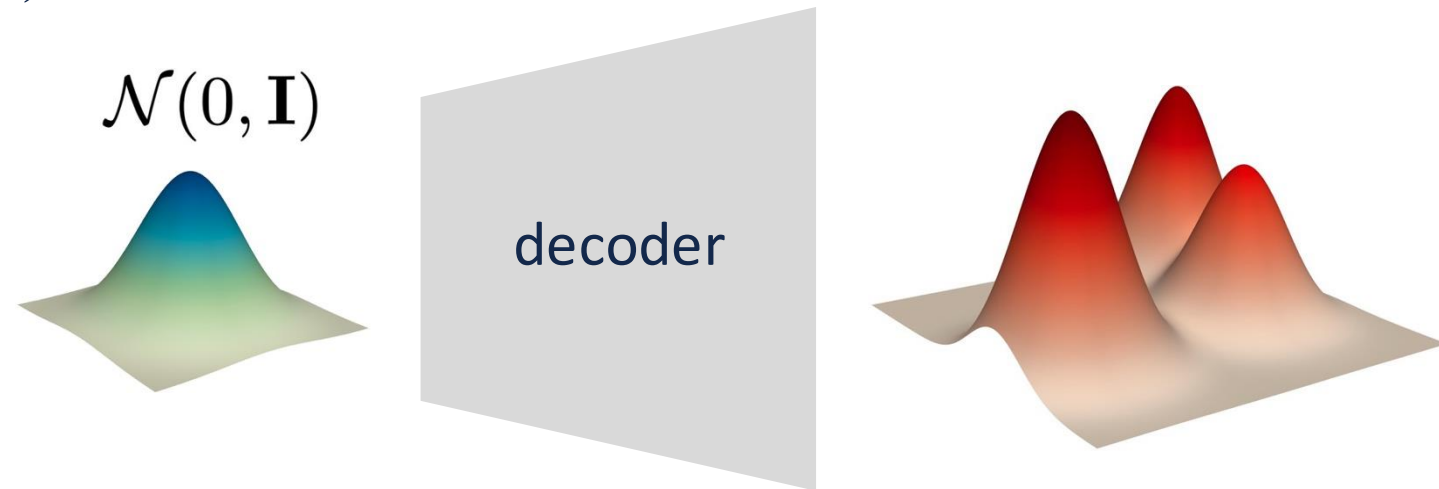


Variational Autoencoder



Inference (generation):

- sample z from: $\mathcal{N}(0, \mathbf{I})$
- map z by decoder net: $g_{\theta}(z)$

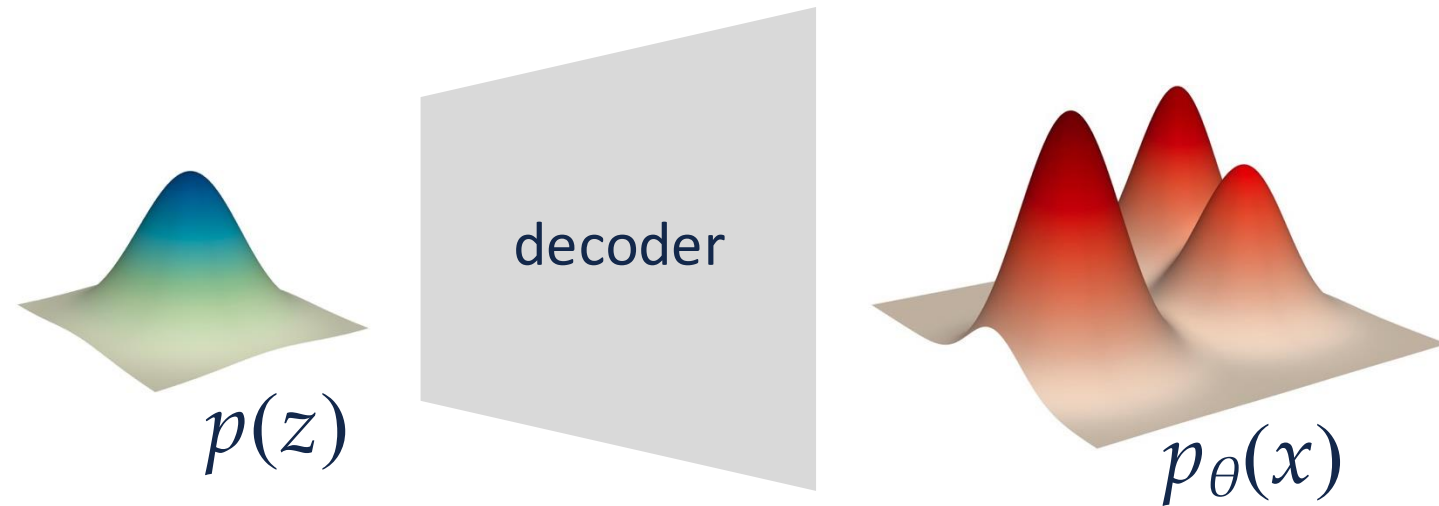


Decoder is a deterministic mapping from one distribution to another.

A view of “Autoencoding Distributions”



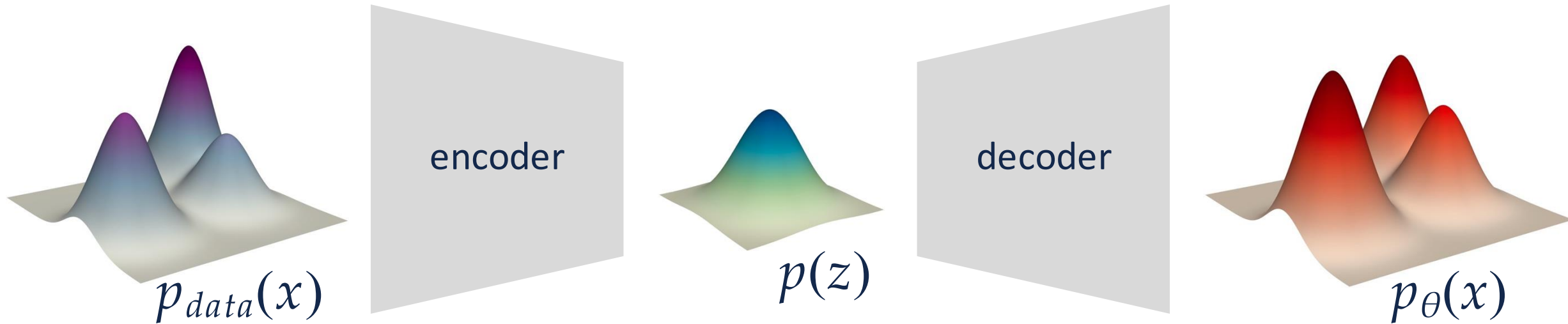
- decoder: maps latent distribution to data distribution



A view of “Autoencoding Distributions”



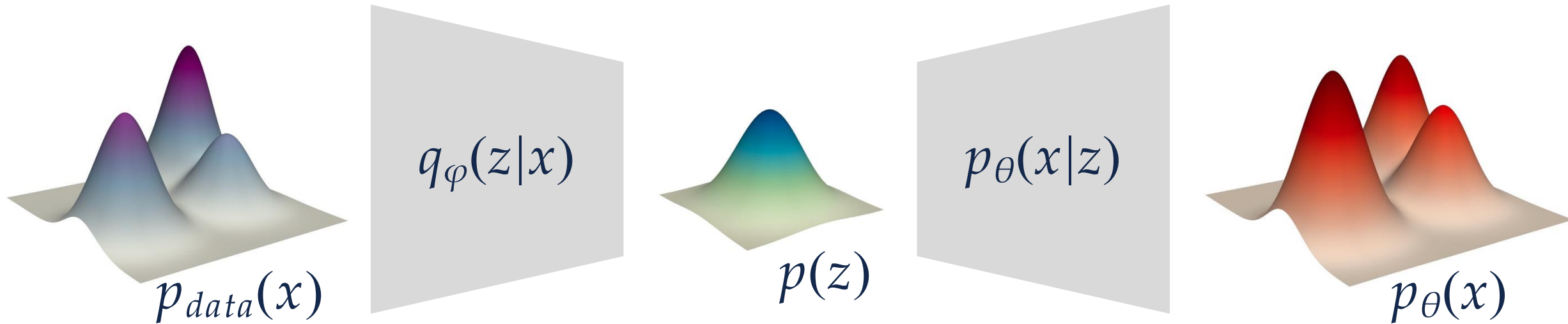
- encoder: maps data distribution to latent distribution
- decoder: maps latent distribution to data distribution



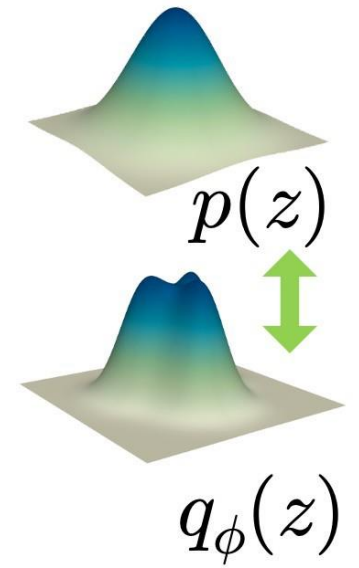
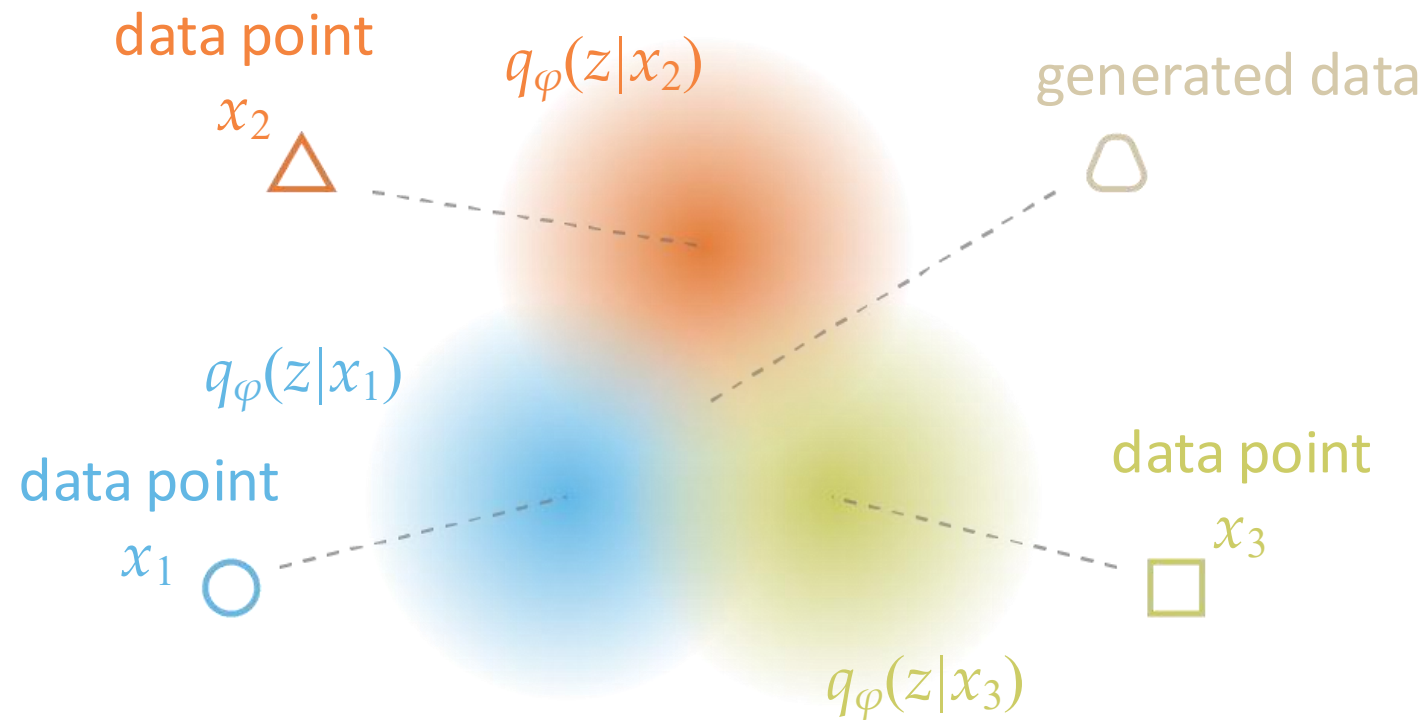
A view of “Autoencoding Distributions”



- encoder: maps data distribution to latent distribution
- decoder: maps latent distribution to data distribution



Illustration

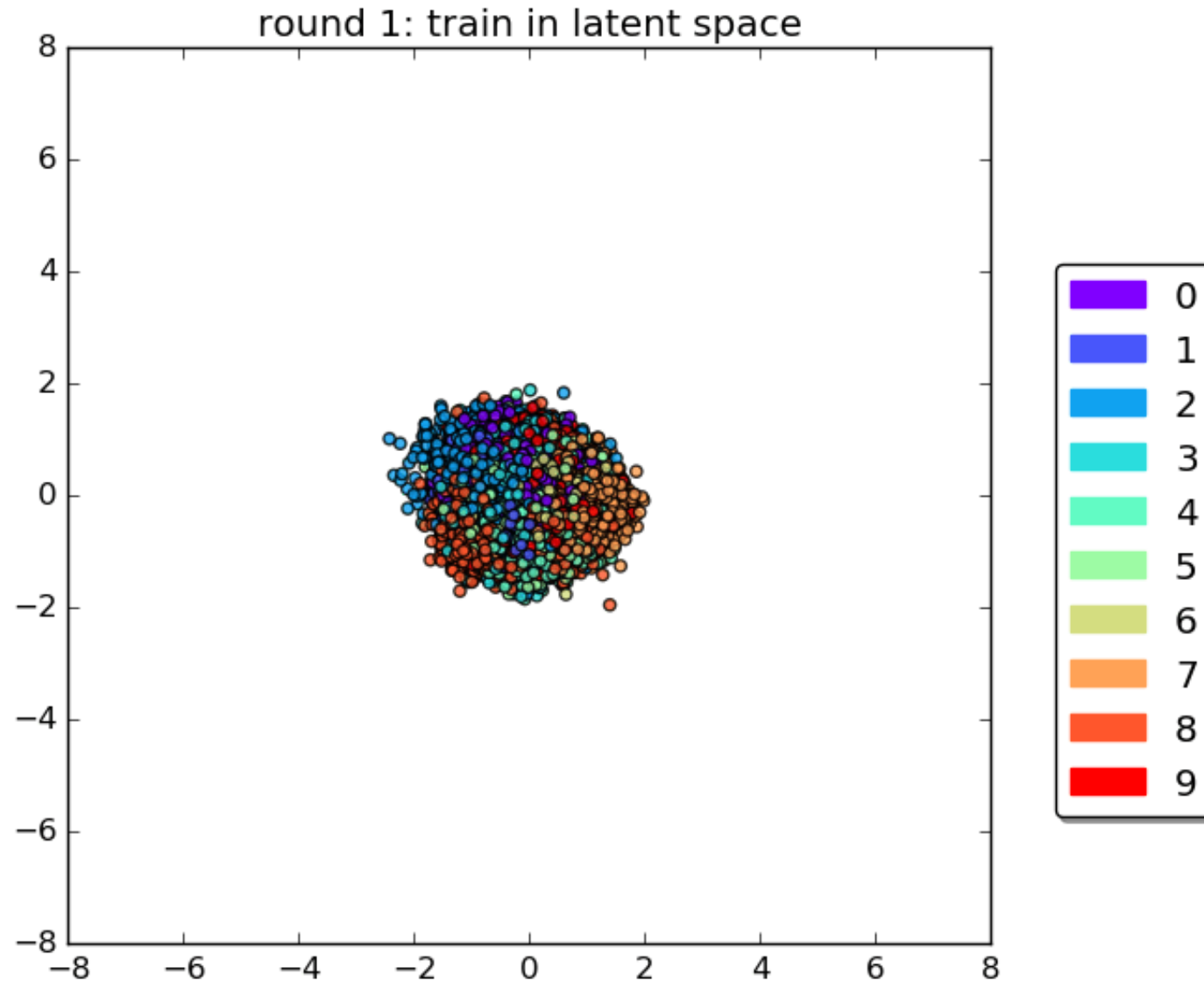


Illustration



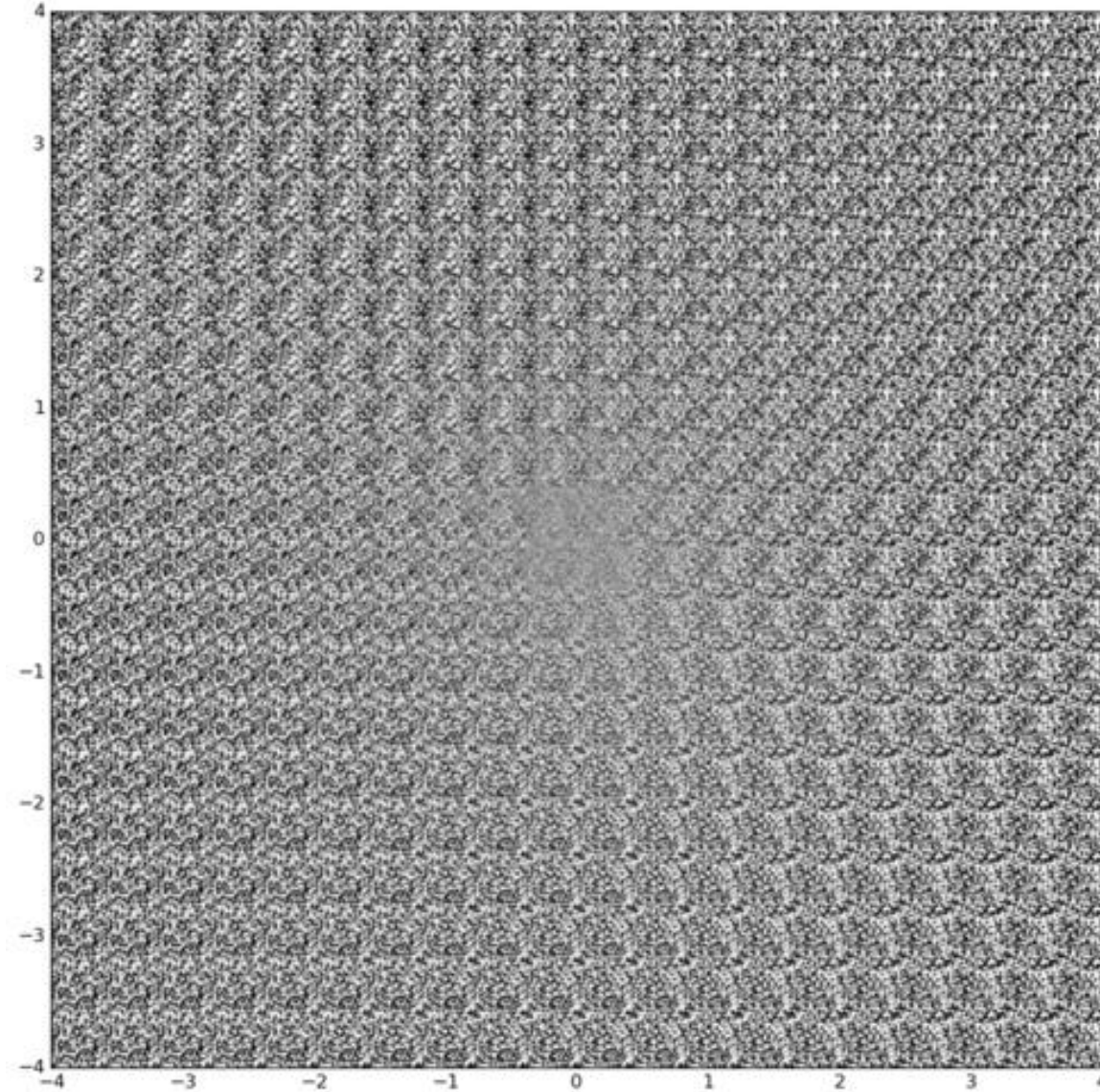
<https://www.youtube.com/watch?v=sV2FOdGqIX0>

VAE: 2D latent space on MNIST



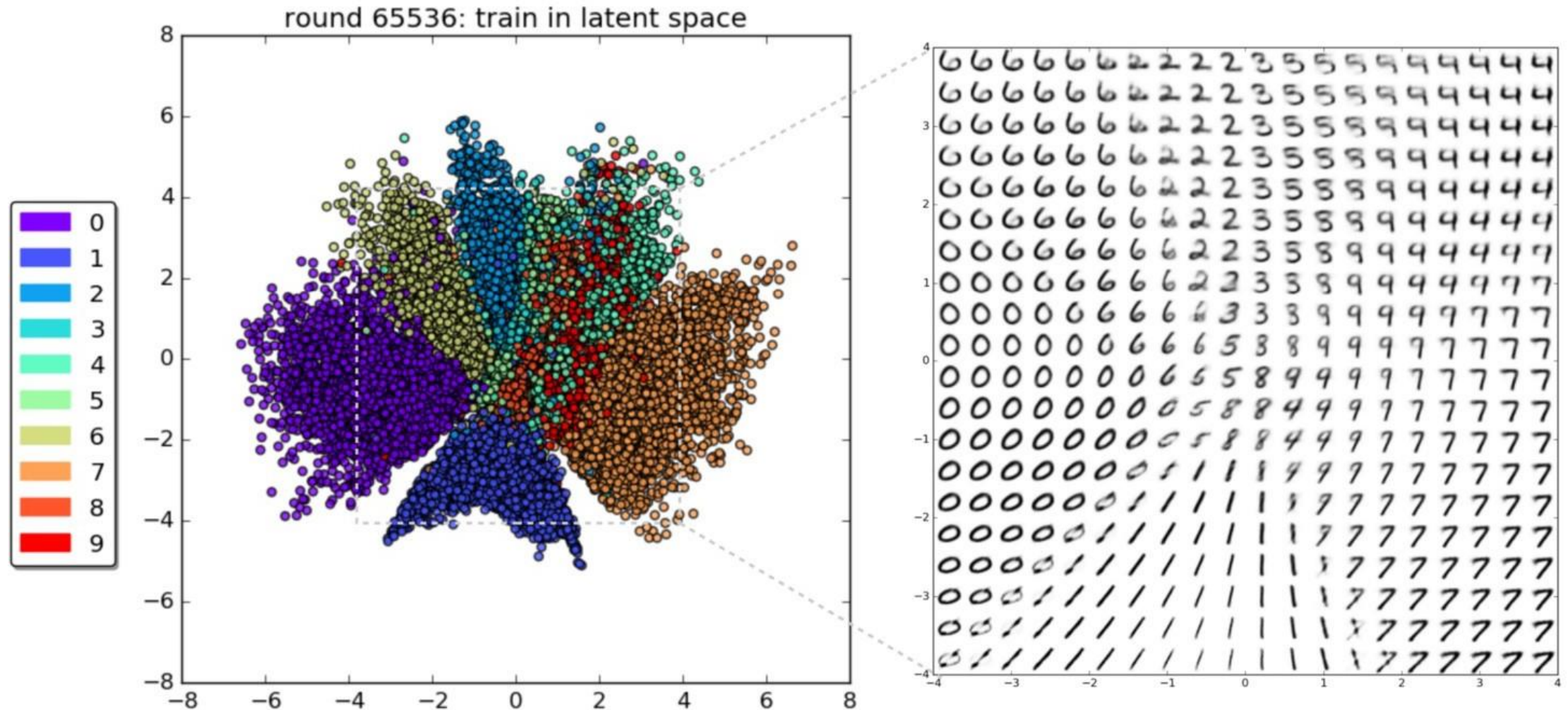
To pdf users: this is animation. Check it on: “Introducing Variational Autoencoders (in Prose and Code)”
<https://blog.fastforwardlabs.com/2016/08/12/introducing-variational-autoencoders-in-prose-and-code.html>

VAE: 2D latent space on MNIST



To pdf users: this is animation. Check it on: “Introducing Variational Autoencoders (in Prose and Code)”
<https://blog.fastforwardlabs.com/2016/08/12/introducing-variational-autoencoders-in-prose-and-code.html>

VAE: 2D latent space on MNIST



“Introducing Variational Autoencoders (in Prose and Code)”

<https://blog.fastforwardlabs.com/2016/08/12/introducing-variational-autoencoders-in-prose-and-code.html>

VAE: 2D latent space on “Frey Face” dataset





The University
of North Carolina
at Chapel Hill