

Generative Models & Embedded Ethics

VAEs, diffusion models, bias in deep learning



CSC420

David Lindell

University of Toronto

cs.toronto.edu/~lindell/teaching/420

Slide credit: FCV, Bill Freeman, Phillip Isola, Tianhong Li

Course Overview



Overview

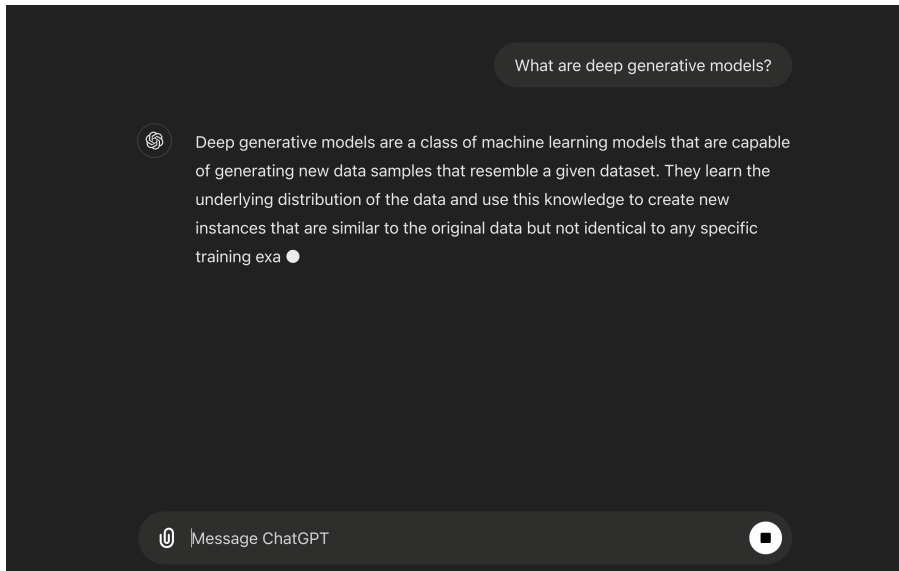
- Introduction
 - What can Generative Models do?
 - What are Generative Models?
 - Generative models and representation learning
- Generative Models: Concepts and Recent Frontiers
 - Variational Autoencoders (VAE)
 - Diffusion models
- Embedded Ethics

Overview

- Introduction
 - What can Generative Models do?
 - What are Generative Models?
 - Generative models and representation learning
- Generative Models: Concepts and Recent Frontiers
 - Variational Autoencoders (VAE)
 - Diffusion models
- Embedded Ethics

The "GenAI" Era

Chatbot and natural language conversation,



The “GenAI” Era

Class-to-image generation, 2019



Generated by BigGAN

The "GenAI" Era

Text-to-image generation, 2022

*'A street sign that reads
"Latent Diffusion" '*

*'A zombie in the
style of Picasso'*

*'A shirt with the inscription:
"I love generative models!" '*



Generated by Latent Diffusion

The "GenAI" Era

Text-to-image generation, 2024



Generated by Stable Diffusion 3 Medium.

Prompt: teddy bear teaching a course, with "generative models" written on blackboard

The "GenAI" Era

Text-to-image generation, 2025

Getting started



Generated by GPT 4o

The "GenAI" Era

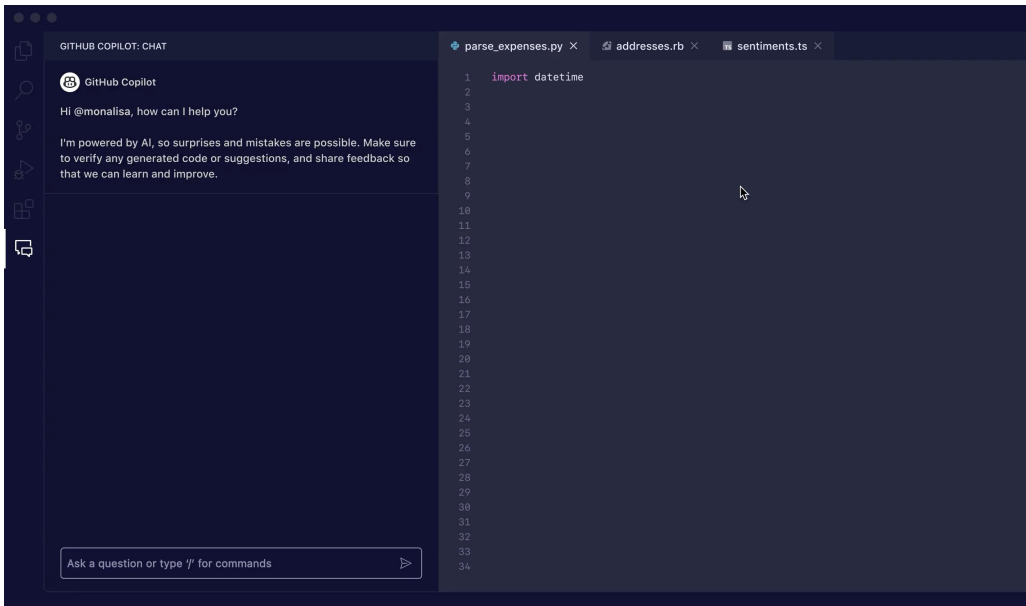
Text-to-video generation, 2024



Generated by Sora

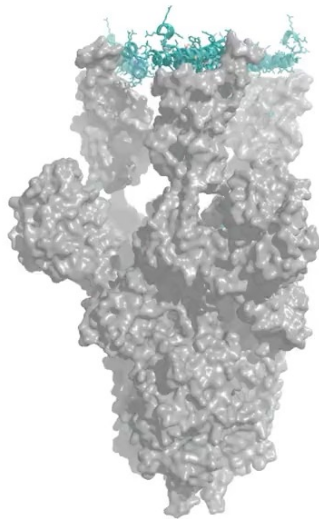
The “GenAI” Era

AI assistant for code generation



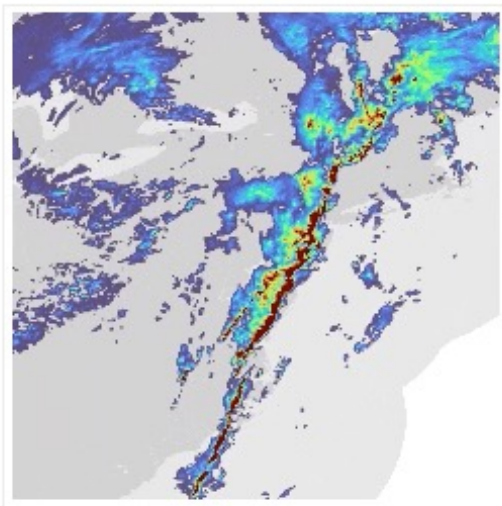
The “GenAI” Era

Protein design and generation

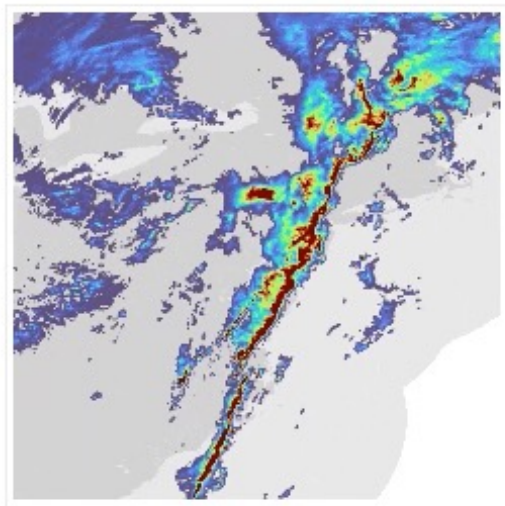


The "GenAI" Era

Weather forecasting



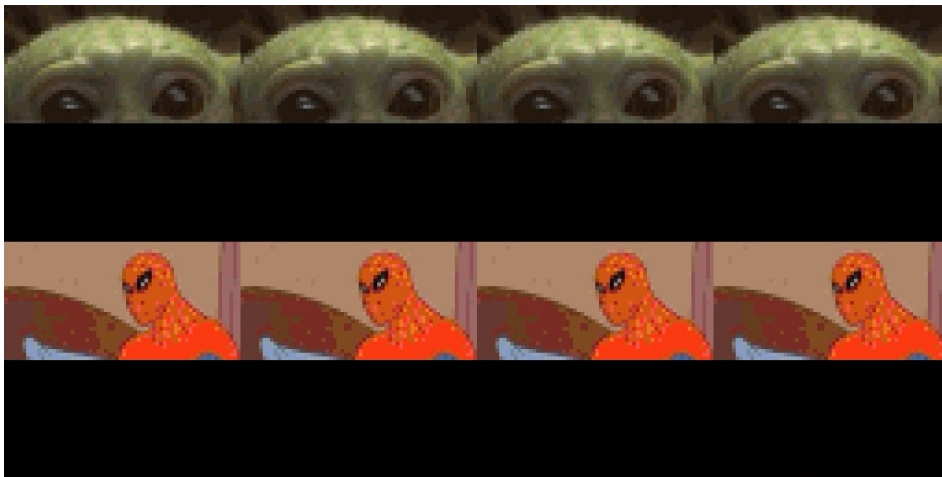
Target



DGMR

Generative Models before the “GenAI” Era

2020, image GPT: Fill bottom part of the image



Generative Models before the “GenAI” Era

2016, Context Encoder: Fill center part of the image



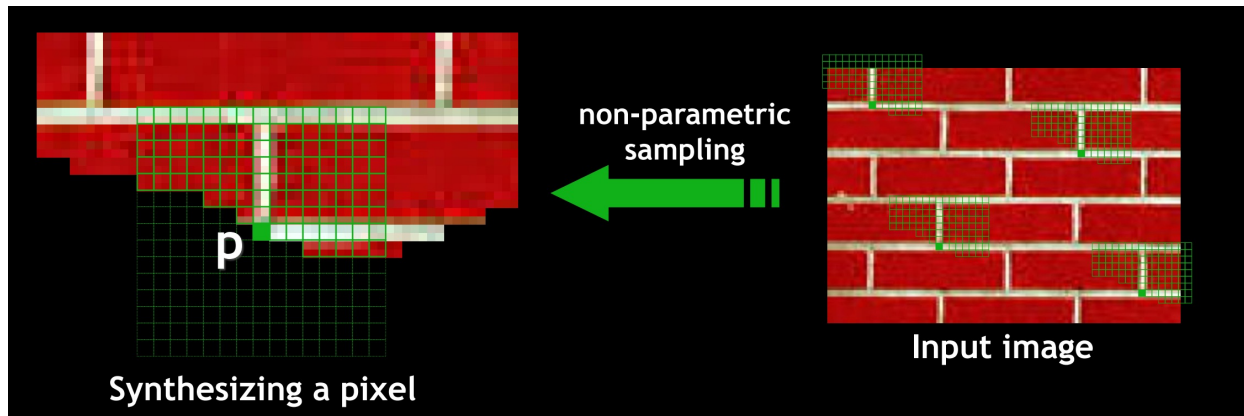
Generative Models before the “GenAI” Era

2009, PatchMatch: Photoshop’s Content-aware Fill



Generative Models before the "GenAI" Era

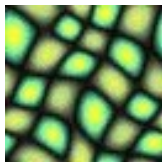
1999, the Efros-Leung algorithm for texture synthesis



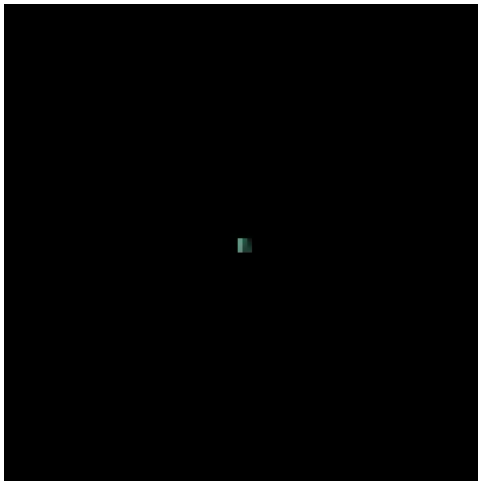
Generative Models before the "GenAI" Era

1999, the Efros-Leung algorithm for texture synthesis

In today's word: this is an **Autoregressive** model

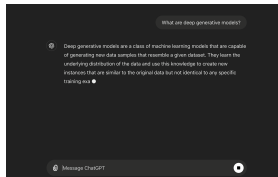


exempla
r



What are Generative Models?

What do these scenarios have in common?



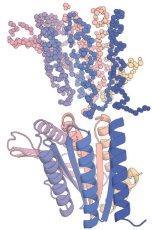
Chatbot



Image generation



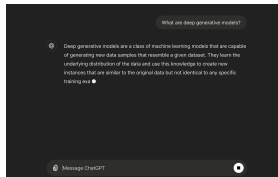
Video generation



Protein generation

What do these scenarios have in common?

- There are **multiple** or infinite predictions to one input.



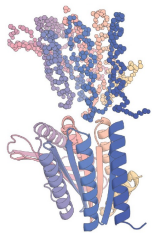
Chatbot



Image generation



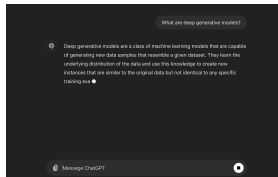
Video generation



Protein generation

What do these scenarios have in common?

- There are **multiple** or infinite predictions to one input.
- Some predictions are more “**plausible**” than some others.



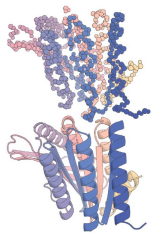
Chatbot



Image generation



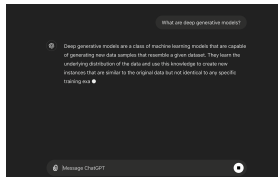
Video generation



Protein generation

What do these scenarios have in common?

- There are **multiple** or infinite predictions to one input.
- Some predictions are more “**plausible**” than some others.
- Training data may contain **no exact solution**.



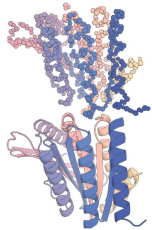
Chatbot



Image generation



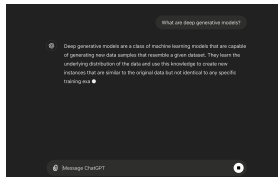
Video generation



Protein generation

What do these scenarios have in common?

- There are **multiple** or infinite predictions to one input.
- Some predictions are more “**plausible**” than some others.
- Training data may contain **no exact solution**.
- Predictions may be **more complex**, more informative, and higher-dimensional than input.



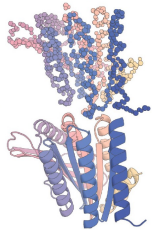
Chatbot



Image generation



Video generation



Protein generation

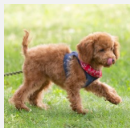
Discriminative vs. Generative models

discriminative

- “sample” $x \Rightarrow$ “label” y
- one desired output

discriminative

x



model



“dog”

y

Discriminative vs. Generative models

discriminative

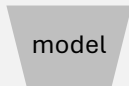
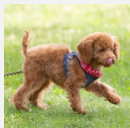
- “sample” $x \Rightarrow$ “label” y
- one desired output

generative

- “label” $y \Rightarrow$ “sample” x
- many possible outputs

discriminative

x



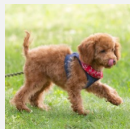
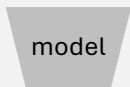
“dog”

y

generative

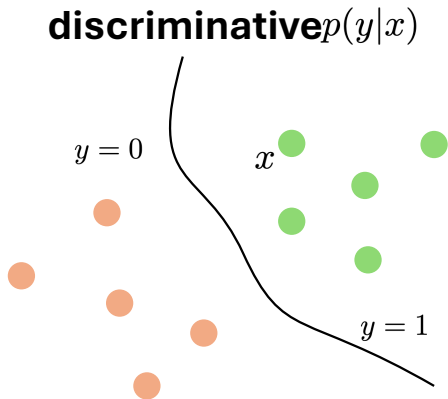
y

“dog”



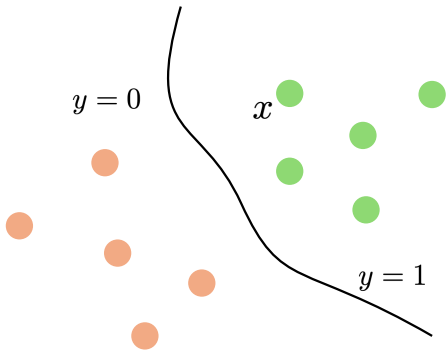
x

Discriminative vs. Generative models

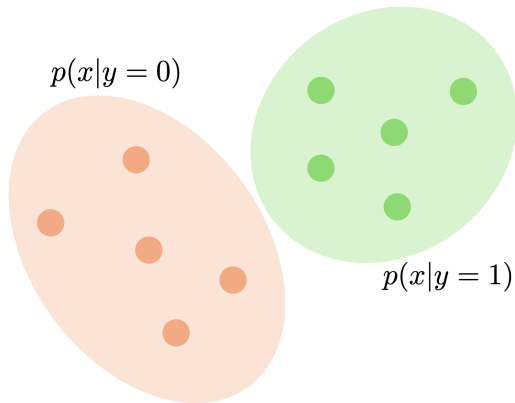


Discriminative vs. Generative models

discriminative $p(y|x)$



generative $p(x|y)$



Formulating Real-world Problems as Generative Models

- Generative models are about $p(x|y)$
- Many problems can be formulated as generative models
- The research of generative models is all about accurately modeling $p(x|y)$

What do we mean by modeling $p(x|y)$

In a **discriminative** model, we care about modeling

$$p(y|x)$$

What do we mean by modeling $p(x|y)$

In a **discriminative** model, we care about modeling

$$p(y|x)$$

this is called the **posterior** — i.e., a probability distribution that represents what we believe about some unknown quantity **after** observing the data (x)

What do we mean by modeling $p(x|y)$

In a **discriminative** model, we care about modeling

$$p(y|x)$$

this is called the **posterior** — i.e., a probability distribution that represents what we believe about some unknown quantity **after** observing the data (x)

typically, we care about finding the ***most likely*** value for y given the data x . This is called maximum a posteriori estimation

$$\operatorname{argmax}_y p(y|x)$$

What do we mean by modeling $p(x|y)$

In a **generative** model, we care about modeling

$$p(x|y)$$

What do we mean by modeling $p(x|y)$

In a **generative** model, we care about modeling

$$p(x|y)$$

this is called the **likelihood** — i.e., the probability that we observe some data (x) given a different set of parameters/labels (y).

What do we mean by modeling $p(x|y)$

In a **generative** model, we care about modeling

$$p(x|y)$$

this is called the **likelihood** — i.e., the probability that we observe some data (x) given a different set of parameters/labels (y).

for a generative model, we typically want to synthesize or sample the data x following $p(x|y)$.

What do we mean by modeling $p(x|y)$

In a **generative** model, we care about modeling

$$p(x|y)$$

this is called the **likelihood** — i.e., the probability that we observe some data (x) given a different set of parameters/labels (y).

- models that can estimate $p(x|y)$ are called “explicit generative models”
- models that can only sample from $p(x|y)$ are called “implicit generative models”

Bayes' Rule

We can model the posterior given the likelihood as

$$\underbrace{p(y \mid x)}_{?} = \frac{p(x \mid y) p(y)}{p(x)}$$

Bayes' Rule

We can model the posterior given the likelihood as

$$\boxed{p(y \mid x)} = \frac{p(x \mid y) p(y)}{p(x)}$$

posterior

Bayes' Rule

We can model the posterior given the likelihood as

$$p(y \mid x) = \frac{p(x \mid y)p(y)}{p(x)}$$

posterior

Bayes' Rule

We can model the posterior given the likelihood as

$$\underset{\text{posterior}}{p(y \mid x)} = \frac{\overset{\text{likelihood}}{p(x \mid y)} p(y)}{p(x)}$$

Bayes' Rule

We can model the posterior given the likelihood as

$$p(y \mid x) = \frac{\overset{\text{likelihood}}{p(x \mid y)} \boxed{p(y)}}{\underset{\text{posterior}}{p(x)}} \quad ?$$

Bayes' Rule

We can model the posterior given the likelihood as

$$\underset{\text{posterior}}{p(y \mid x)} = \frac{\overset{\text{likelihood}}{p(x \mid y)} \boxed{p(y)}}{\quad} \text{prior}$$

a prior could encode, e.g., how common each class label is

Bayes' Rule

We can model the posterior given the likelihood as

$$\underset{\text{posterior}}{p(y \mid x)} = \frac{\overset{\text{likelihood}}{p(x \mid y)} \underset{\text{prior}}{p(y)}}{\underset{?}{p(x)}}$$

Bayes' Rule

We can model the posterior given the likelihood as

$$p(y | x) = \frac{\overset{\text{likelihood}}{p(x | y)} \overset{\text{prior}}{p(y)}}{\underset{\text{posterior}}{\boxed{p(x)}} \underset{\text{marginal likelihood/data distribution}}{}}$$

$$p(y | x) = \frac{p(x | y) p(y)}{\boxed{\int p(x | y) p(y) dy}}$$

marginal likelihood/
data distribution

Bayes' Rule

We can model the posterior given the likelihood as

$$p(y \mid x) = \frac{\overset{\text{likelihood}}{p(x \mid y)} \overset{\text{prior}}{p(y)}}{\underset{\text{posterior}}{p(x)} \underset{\text{marginal likelihood/data distribution}}{p(x)}}$$

can explicit generative models be made discriminative?

Why is it hard?

- Explicit generative models can be discriminative: Bayes' rule

$$\underset{\text{discriminative}}{p(y|x)} = \underset{\text{generative}}{p(x|y)} \frac{p(y)}{p(x)}$$

assuming known prior on label y

constant for given x

Why is it hard?

- Explicit generative models can be discriminative: Bayes' rule

$$\underset{\text{discriminative}}{p(y|x)} = \underset{\text{generative}}{p(x|y)} \frac{p(y)}{p(x)}$$

assuming known prior on label y

constant for given x

- Can discriminative models be generative?

Why is it hard?

- Explicit generative models can be discriminative: Bayes' rule

$$\underset{\text{discriminative}}{p(y|x)} = \underset{\text{generative}}{p(x|y)} \frac{p(y)}{p(x)}$$

assuming known prior on label y

constant for given x

- Can discriminative models be generative?

$$\underset{\text{generative}}{p(x|y)} = \underset{\text{discriminative}}{p(y|x)} \frac{p(x)}{p(y)}$$

still need to model prior distribution of data x

constant for given y

Why is it hard?

- Explicit generative models can be discriminative: Bayes' rule

$$\underset{\text{discriminative}}{p(y|x)} = \underset{\text{generative}}{p(x|y)} \frac{p(y)}{p(x)}$$

assuming known prior on label y

constant for given x

- Can discriminative models be generative?

$$\underset{\text{generative}}{p(x|y)} = \underset{\text{discriminative}}{p(y|x)} \frac{p(x)}{p(y)}$$

still need to model prior distribution of data x

constant for given y

- The challenge is about **modeling data distributions**

Formulating Real-world Problems as Generative Modeling

Formulating Real-world Problems as Generative Models

- Generative models perform (or describe) the synthesis of data
- In other words, generative models are about $p(x|y)$

What can be y ?

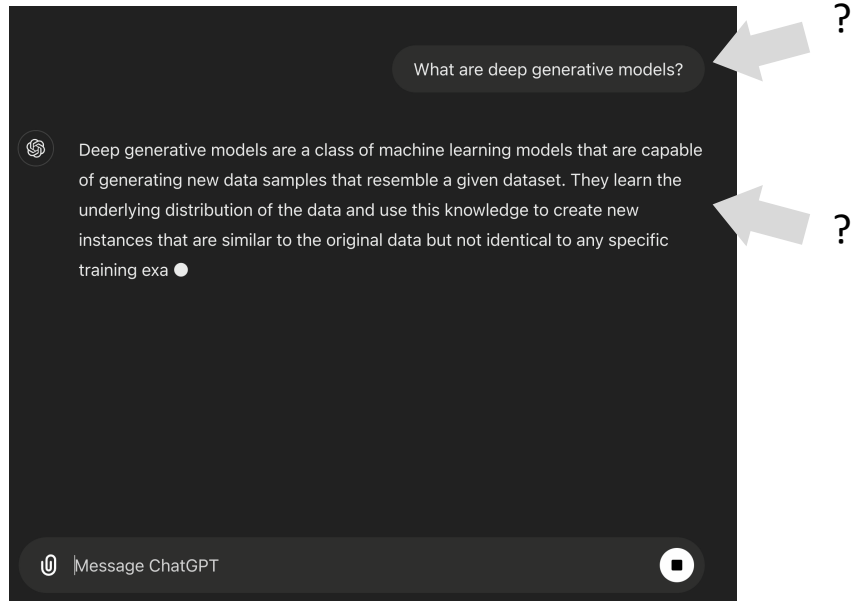
- condition
- constraint
- labels
- attributes
- ...
- more abstract
- less informative

What can be x ?

- “data”
- samples
- observations
- measurements
- ...
- more concrete
- more informative

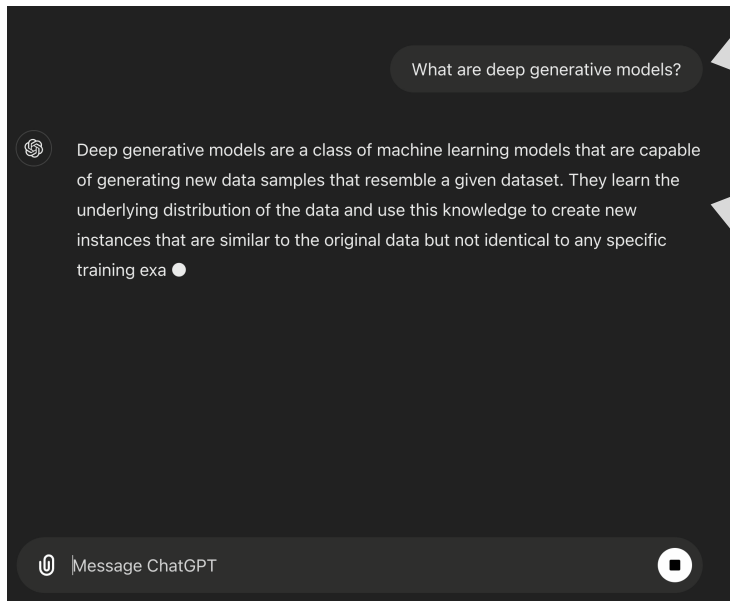
Case study: Formulating problems as $p(x|y)$

- **Natural language conversation**



Case study: Formulating problems as $p(x|y)$

- **Natural language conversation**



y: prompt

x: response of the chatbot

Case study: Formulating problems as $p(x|y)$

- **Text-to-image/video generation**

*Prompt: teddy bear teaching a course, with
"generative models" written on blackboard*



Case study: Formulating problems as $p(x|y)$

- **Text-to-image/video generation**

*Prompt: teddy bear teaching a course, with
"generative models" written on blackboard*



y: text prompt



x: generated visual content

Case study: Formulating problems as $p(x|y)$

- **Text-to-3D structure generation**



“motorcycle”



“mech suit”



“ghost lantern”



“furry fox head”



?



?



“dresser”



“swivel chair”



“astronaut”



“mushroom house”

Case study: Formulating problems as $p(x|y)$

- **Text-to-3D structure generation**



“motorcycle”



“mech suit”



“ghost lantern”



“furry fox head”



x: generated
3D structures



y: text prompt



“dresser”



“swivel chair”



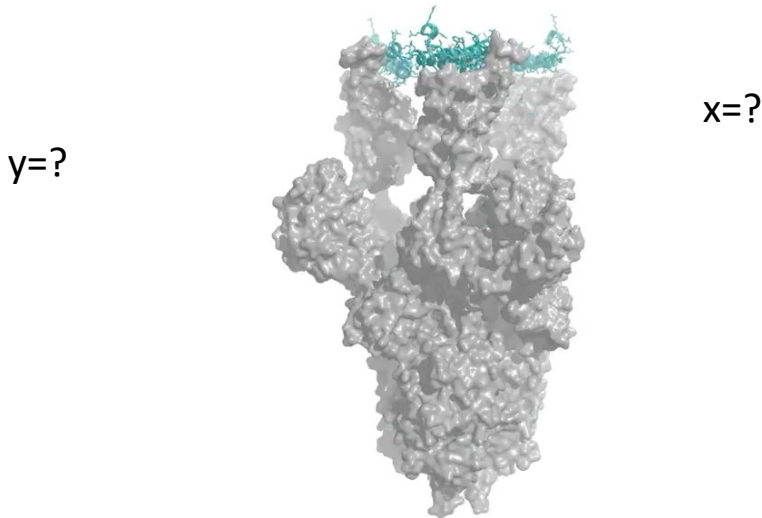
“astronaut”



“mushroom house”

Case study: Formulating problems as $p(x|y)$

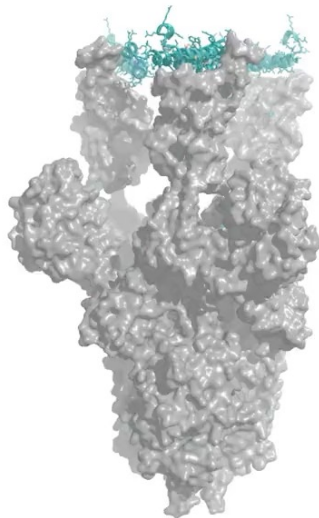
- **Protein structure generation**



Case study: Formulating problems as $p(x|y)$

- **Protein structure generation**

y:
condition/constraint
(e.g., symmetry)

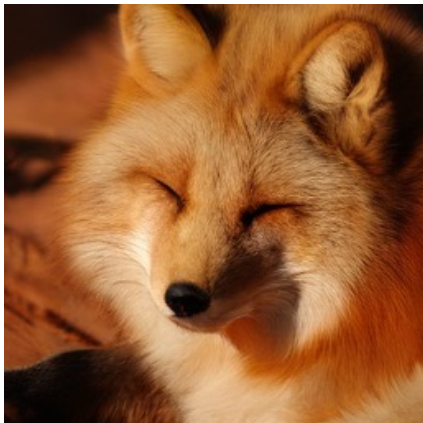


x: generated
protein
structures

Case study: Formulating problems as $p(x|y)$

- **Class-conditional image generation**

"red fox"



?

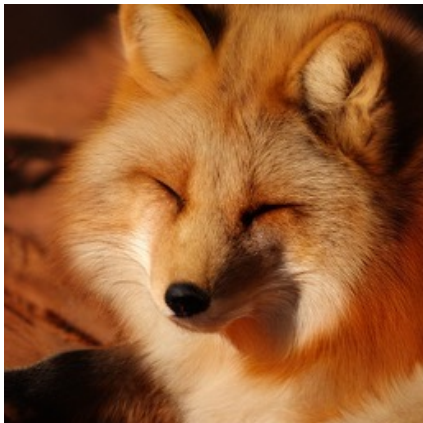


?

Case study: Formulating problems as $p(x|y)$

- **Class-conditional image generation**

"red fox"



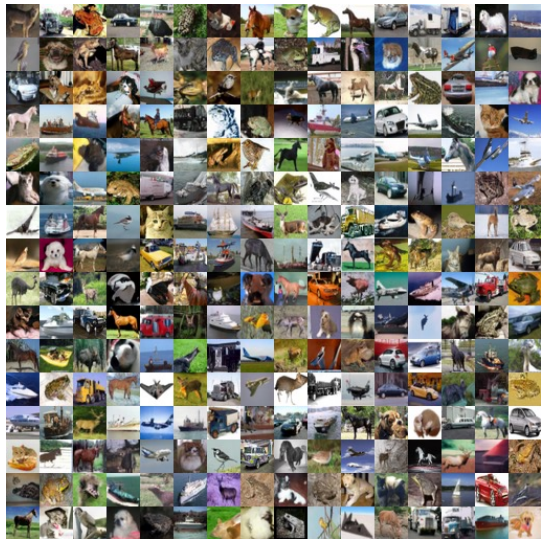
y: class label



x: generated image

Case study: Formulating problems as $p(x|y)$

- “Unconditional” image generation

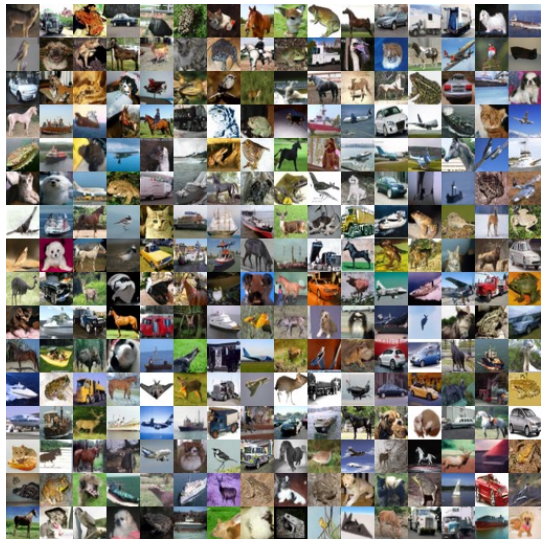


$y: ?$

$x ?$

Case study: Formulating problems as $p(x|y)$

- “Unconditional” image generation



y : an implicit condition
“images following CIFAR10 distribution”

x : generated CIFAR10-like images

- $p(x|y)$: images $\sim$ CIFAR10
- $p(x)$: all images

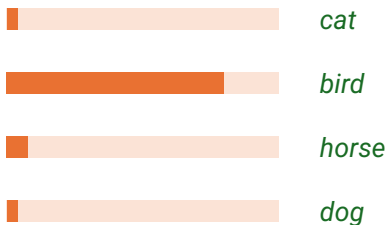
Case study: Formulating problems as $p(x|y)$

- **Classification** (a generative perspective)

y : ?



x : ?



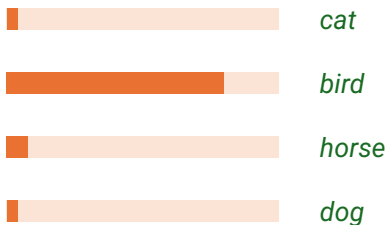
Case study: Formulating problems as $p(x|y)$

- **Classification** (a generative perspective)

y: an image as the “condition”



x: probability of classes
conditioned on the
image

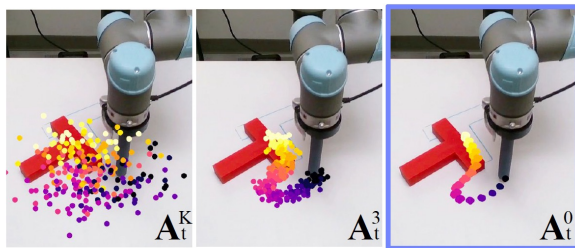
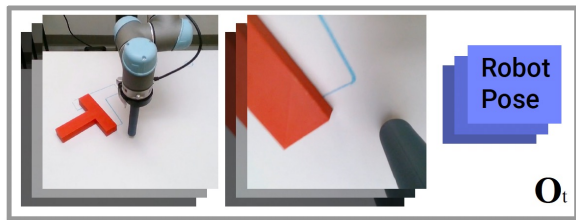


Case study: Formulating problems as $p(x|y)$

- **Policy Learning in Robotics**

$x: ?$

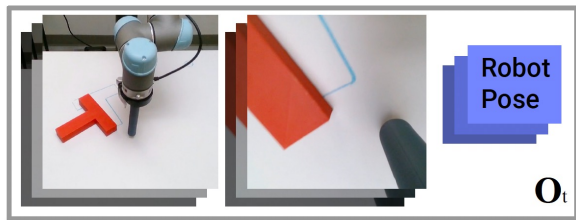
$y: ?$



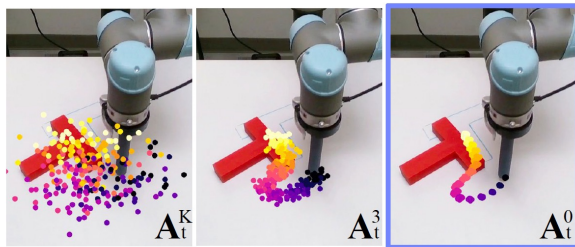
Case study: Formulating problems as $p(x|y)$

- **Policy Learning in Robotics**

y : visual and other
sensory observations



x : policies
(probability of actions)



Generative Models and Representation Learning

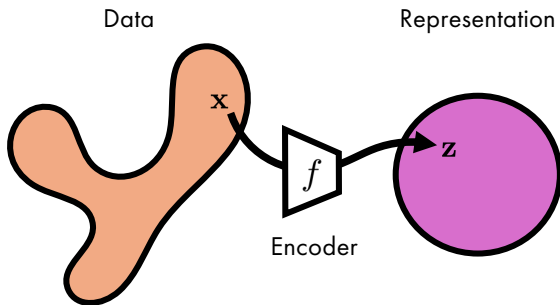
Generative Models and Representation Learning

- **Representation learning:**

- Learning to **represent data instances**

- map data to feature: $x \rightarrow f(x)$

- minimize objective: $\min_f \text{Loss}(f(x))$



Generative Models and Representation Learning

- **Representation learning:**

- Learning to **represent data instances**

- map data to feature: $x \rightarrow f(x)$

- minimize objective: $\min_f \text{Loss}(f(x))$

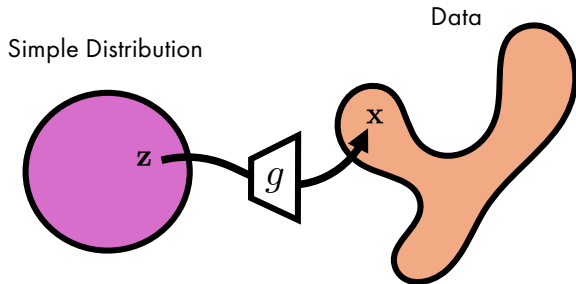
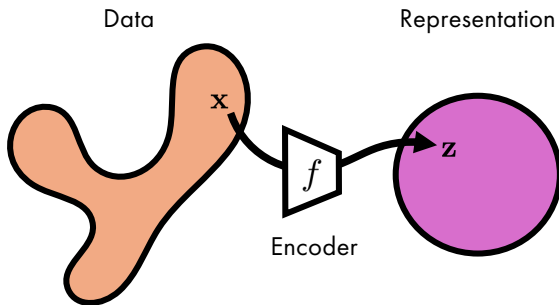
- **Generative models:**

- Learning to **represent probability distributions**

- map a simple distribution to a complex one

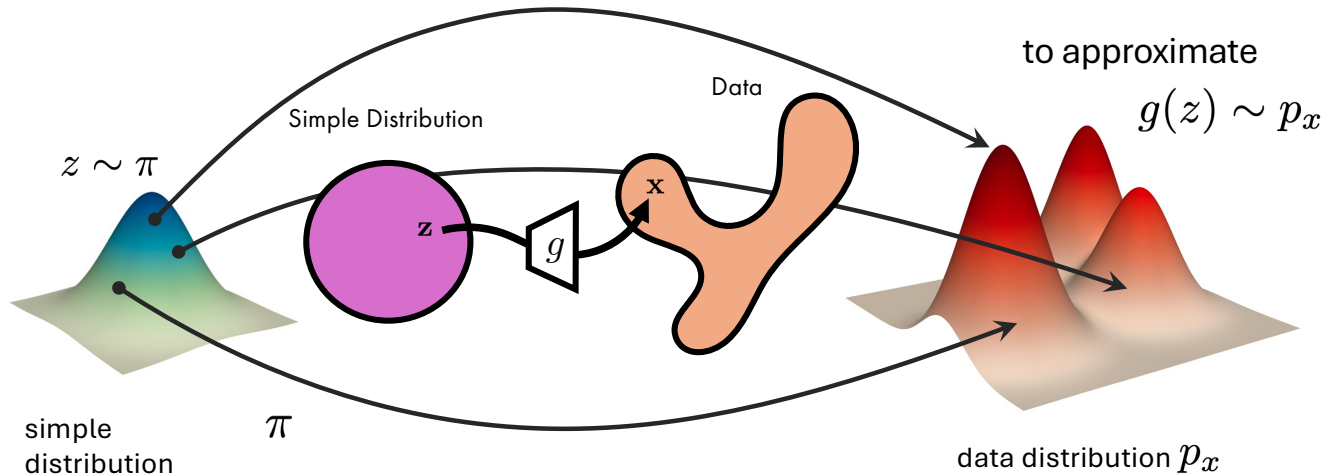
- minimize loss, so that $p(g(z)) \sim p(x)$:

$$\min_g \text{Loss}(x, g(z))$$



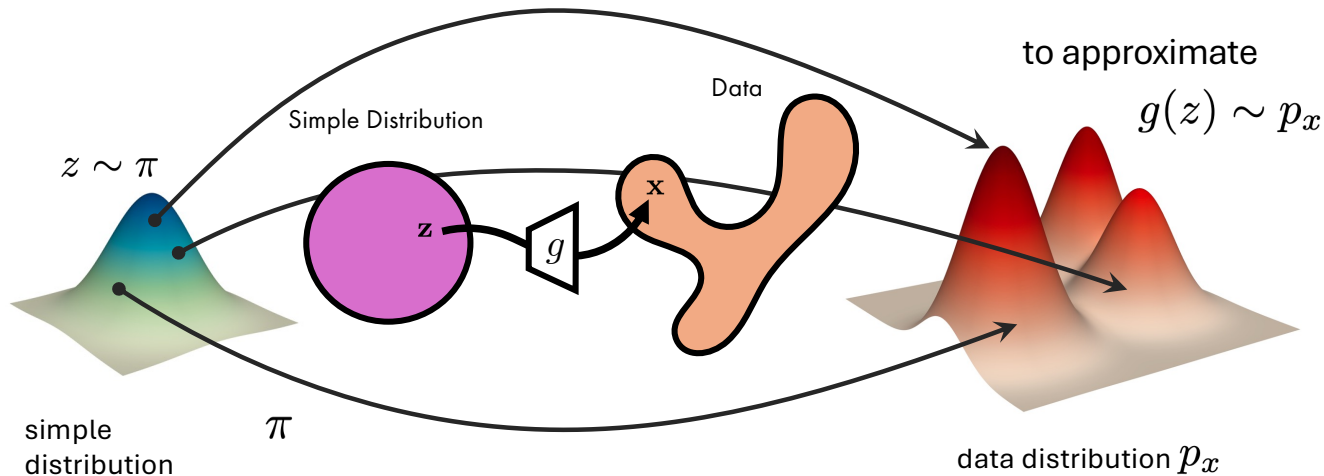
Learning to represent probability distributions

- How to sample complex data distributions $p(x)$?
- Learn the **mapping** from simple to complex distributions!



Learning to represent probability distributions

- How do we “learn” this mapping? What is the loss?

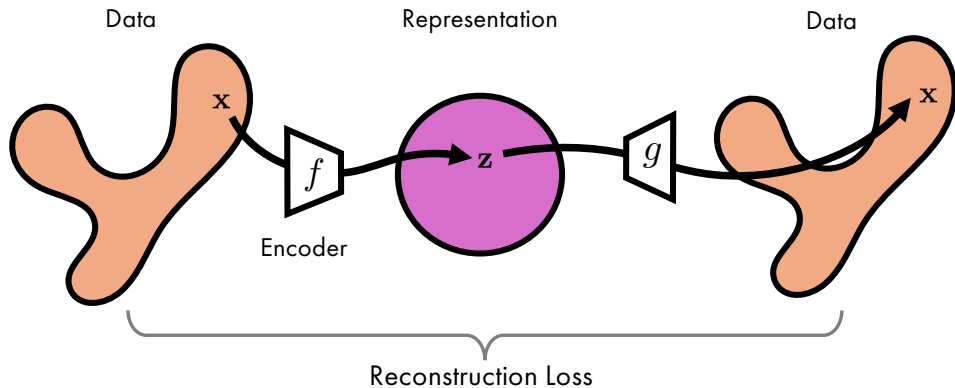


Overview

- Introduction
 - What can Generative Models do?
 - What are Generative Models?
 - Generative models and representation learning
- Generative Models: Concepts and Recent Frontiers
 - Variational Autoencoders (VAE)
 - Diffusion models
- Embedded Ethics

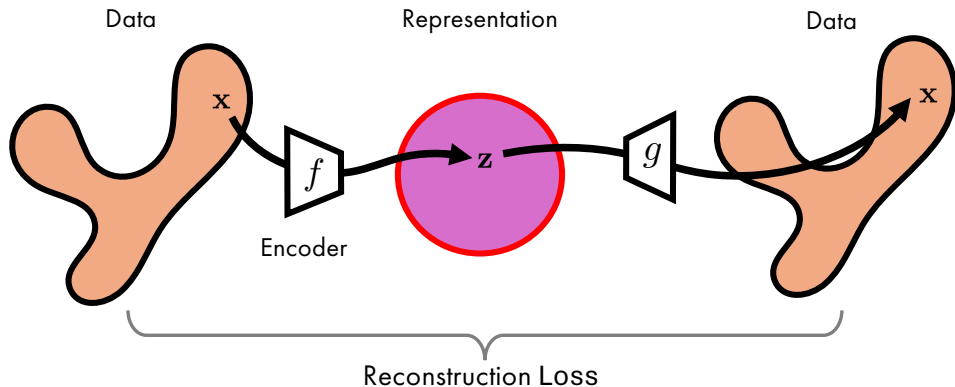
Autoencoder

- Recall Autoencoder in representation learning
- We have a compact “representation space”
- We have a decoder that maps a representation to data



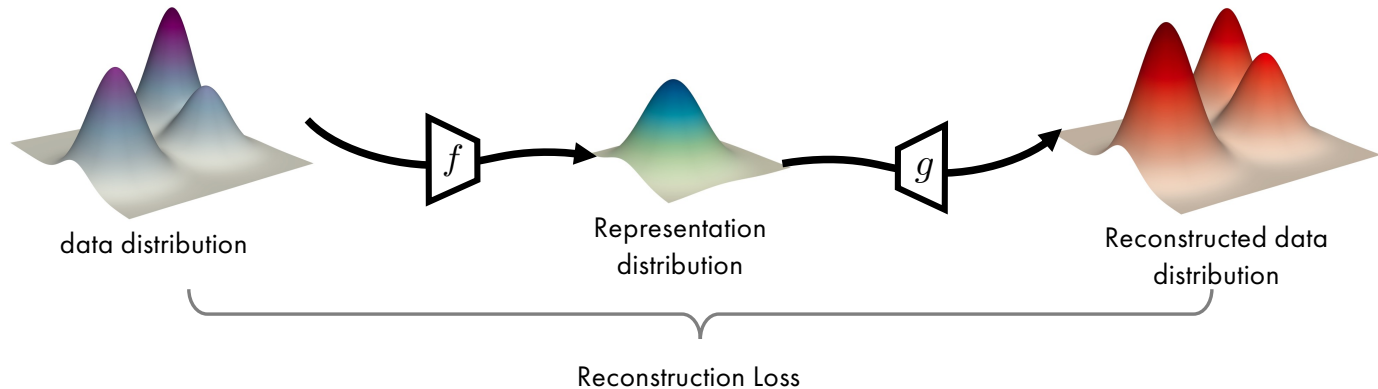
Autoencoder

- Recall Autoencoder in representation learning
- We have a compact “representation space”
- We have a decoder that maps a representation to data
- Problem: we still cannot sample this representation



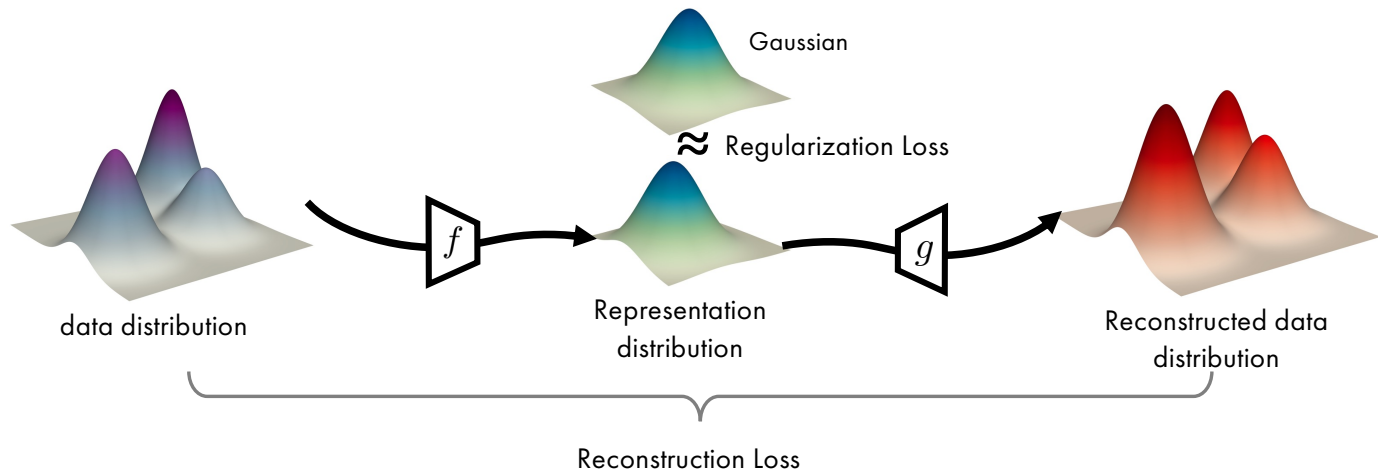
Autoencoder

- Recall Autoencoder in representation learning
- We have a compact “representation space”
- We have a decoder that maps a representation to data
- Problem: we still cannot sample this representation



Variational Autoencoder

- Almost the same as Autoencoder
- Need to regularize the representation distribution, so that it follows a “samplable” distribution, e.g., Gaussian.



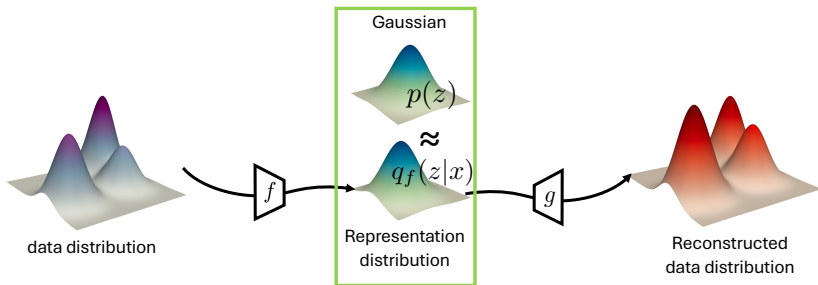
Variational Autoencoder

Regularization loss

$$\mathcal{D}_{\text{KL}}(q_f(z|x) || p(z))$$

Example: Gaussian prior

- let $p(z) = \mathcal{N}(z | 0, \mathbf{I})$
- model $q_f(z|x)$ by Gaussian: $\mathcal{N}(z | \mu, \sigma)$



Variational Autoencoder

Regularization loss

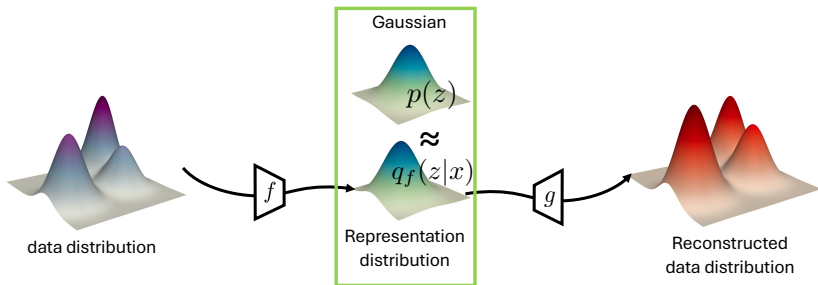
$$\mathcal{D}_{\text{KL}}(q_f(z|x) || p(z))$$

Example: Gaussian prior

- let $p(z) = \mathcal{N}(z | 0, \mathbf{I})$
- model $q_f(z|x)$ by Gaussian: $\mathcal{N}(z | \mu, \sigma)$
- map x by encoder net:

$$f_\phi(x) \rightarrow \mu, \sigma$$

network estimates
distribution's
parameters



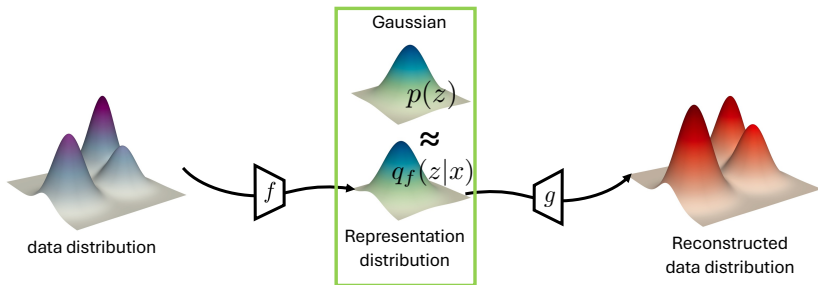
Variational Autoencoder

Regularization loss

$$\mathcal{D}_{\text{KL}}(q_f(z|x) || p(z))$$

Example: Gaussian prior

- let $p(z) = \mathcal{N}(z | 0, \mathbf{I})$
- model $q_f(z|x)$ by Gaussian: $\mathcal{N}(z | \mu, \sigma)$
- map x by encoder net:
- compute loss analytically: $\mathcal{D}_{\text{KL}}(\mathcal{N}(z | \mu, \sigma) || \mathcal{N}(z | 0, \mathbf{I}))$



$$f_\phi(x) \rightarrow \mu, \sigma$$

network estimates
distribution's
parameters

Variational Autoencoder

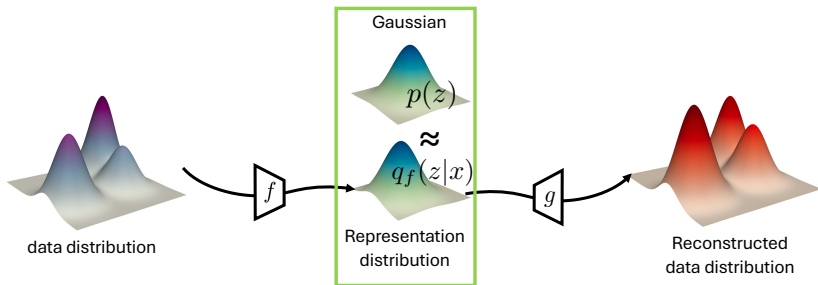
Regularization loss

$$\mathcal{D}_{\text{KL}}(q_f(z|x) || p(z))$$

Example: Gaussian prior

- let $p(z) = \mathcal{N}(z | 0, \mathbf{I})$
- model $q_f(z|x)$ by Gaussian: $\mathcal{N}(z | \mu, \sigma)$
- map x by encoder net:
- compute loss analytically: $\mathcal{D}_{\text{KL}}(\mathcal{N}(z | \mu, \sigma) || \mathcal{N}(z | 0, \mathbf{I}))$
- fixed variance $\sigma = 1 \Rightarrow \text{L2 loss on } \mu$

$$\mathcal{D}_{\text{KL}}(q_f(z|x) || p(z)) = ||\mu||^2$$



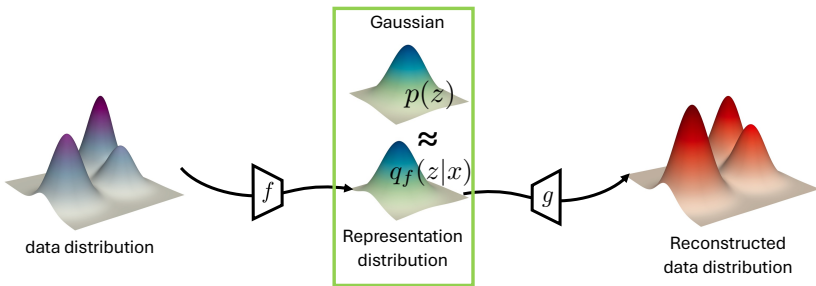
$$f_\phi(x) \rightarrow \mu, \sigma$$

network estimates
distribution's
parameters

Variational Autoencoder

Reconstruction loss

$$-\mathbb{E}_{q_f(z|x)} [\log p_\theta(x | z)]$$

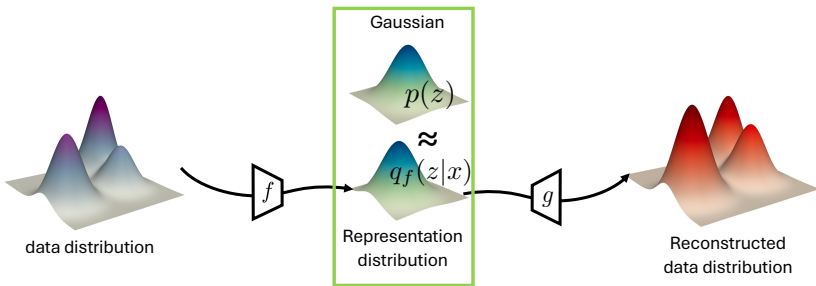


Variational Autoencoder

Reconstruction loss

$$-\mathbb{E}_{q_f(z|x)} [\log p_\theta(x | z)]$$

Example: Assume $p(x|z)$
to be Gaussian



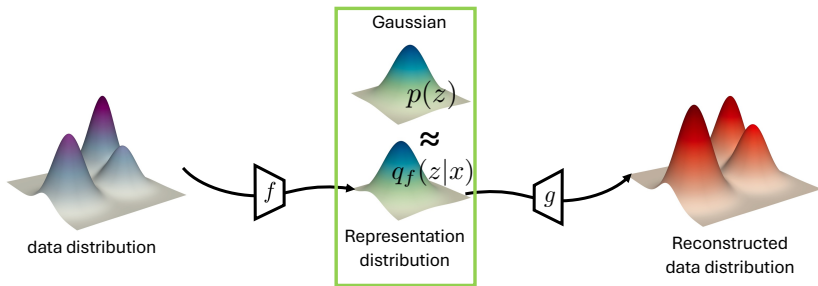
Variational Autoencoder

Reconstruction loss

$$-\mathbb{E}_{q_f(z|x)} [\log p_\theta(x | z)]$$

Example: Assume $p(x|z)$
to be Gaussian

- we approximate the expectation with a single sample drawn from $q_f(z|x)$



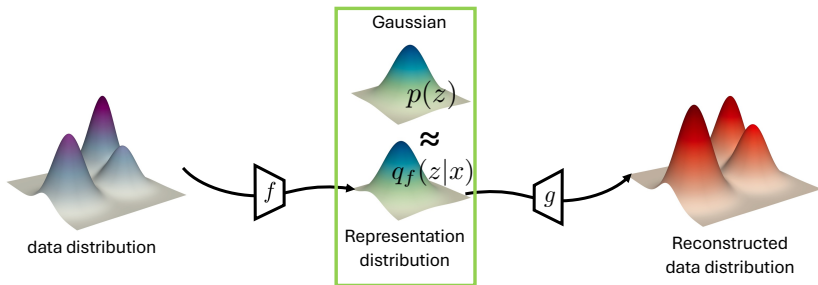
Variational Autoencoder

Reconstruction loss

$$-\mathbb{E}_{q_f(z|x)} [\log p_\theta(x | z)]$$

Example: Assume $p(x|z)$
to be Gaussian

- we approximate the expectation with a single sample drawn from $q_f(z|x)$
- model each pixel as independent with its own mean and constant variance



Variational Autoencoder

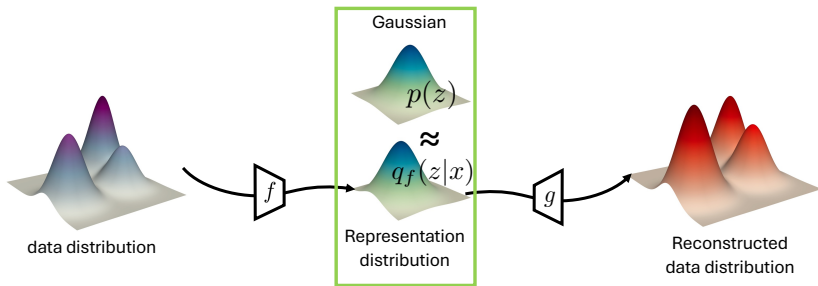
Reconstruction loss

$$-\mathbb{E}_{q_f(z|x)} [\log p_\theta(x | z)]$$

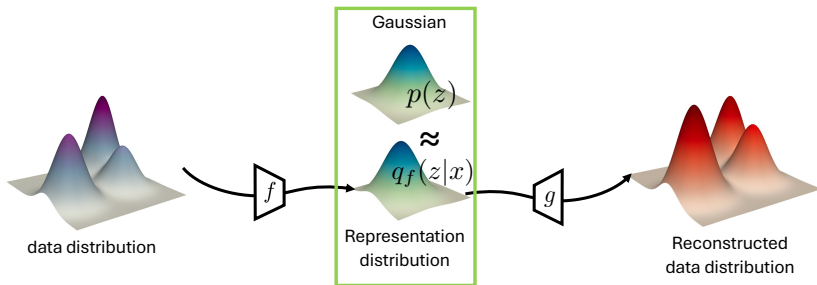
Example: Assume $p(x|z)$ to be Gaussian

- we approximate the expectation with a single sample drawn from $q_f(z|x)$
- model each pixel as independent with its own mean and constant variance
- then the reconstruction loss reduces to simple L2 loss

$$\|x - \mu(z)\|_2^2$$



Variational Autoencoder



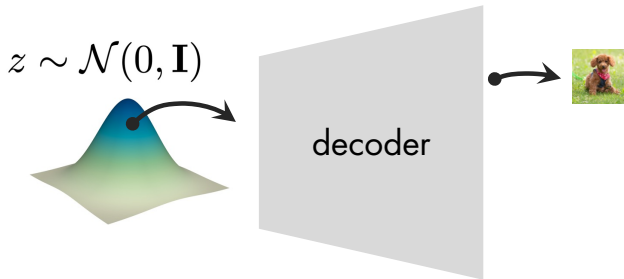
$$\mathcal{L}(g, f; \mathbf{x}) = \underbrace{-\mathbb{E}_{q_f(z|\mathbf{x})} [\log p_g(\mathbf{x} | z)]}_{\text{reconstruction term}} + \underbrace{D_{KL}(q_f(z | \mathbf{x}) || p(z))}_{\text{regularization term}}$$

you'll derive and implement in your homework!

Variational Autoencoder

Inference (generation):

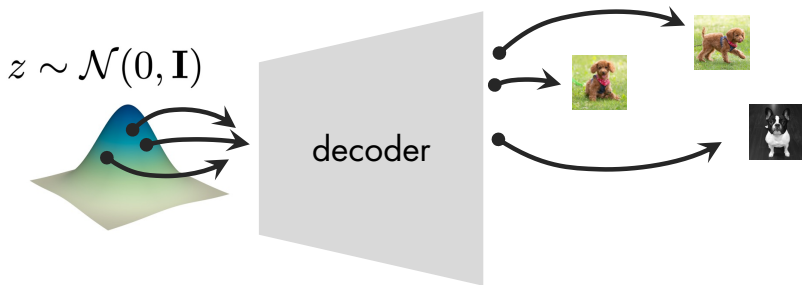
- sample z from: $\mathcal{N}(0, \mathbf{I})$
- map z by decoder net



Variational Autoencoder

Inference (generation):

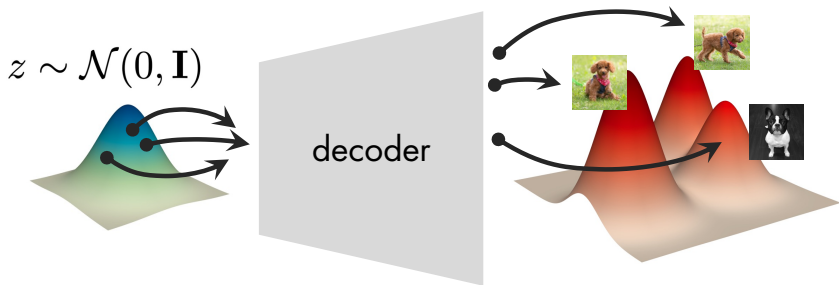
- sample z from: $\mathcal{N}(0, \mathbf{I})$
- map z by decoder net



Variational Autoencoder

Inference (generation):

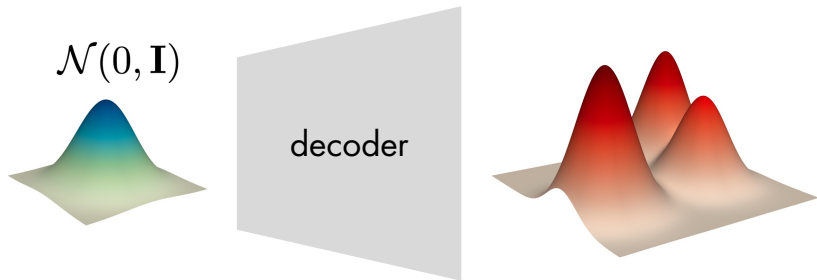
- sample z from: $\mathcal{N}(0, \mathbf{I})$
- map z by decoder net



Variational Autoencoder

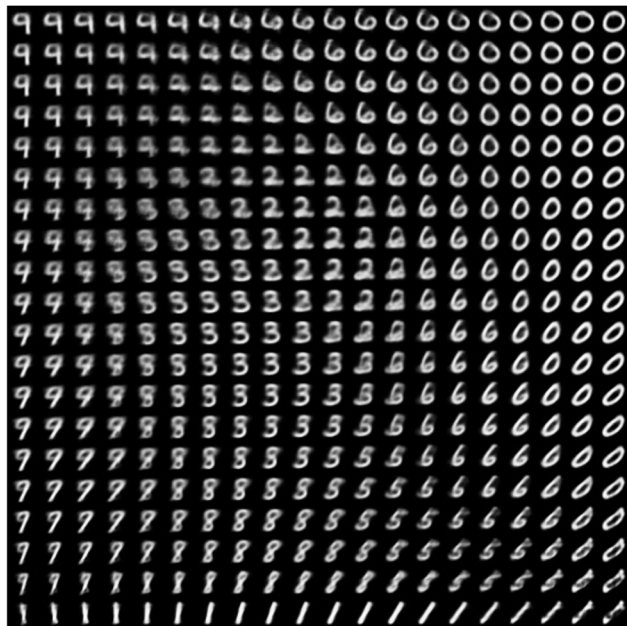
Inference (generation):

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



Decoder is a deterministic mapping from one distribution to another.

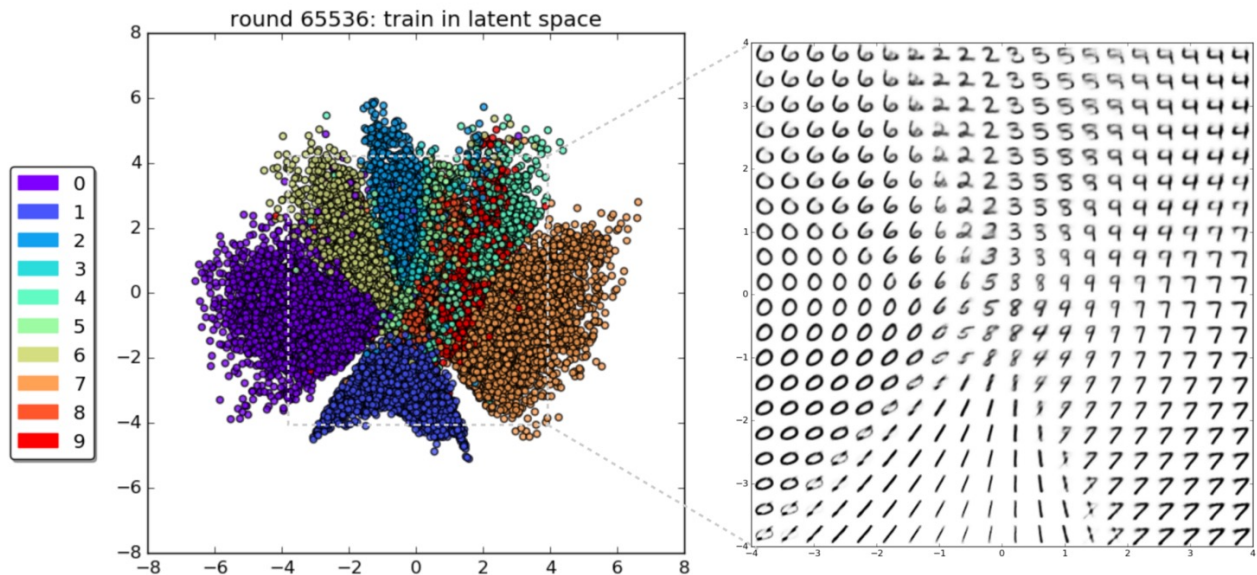
Example: 2D latent space on MNIST



“Convolutional Variational Autoencoder”

<https://colab.research.google.com/github/tensorflow/docs/blob/master/site/en/tutorials/generative/cvae.ipynb>

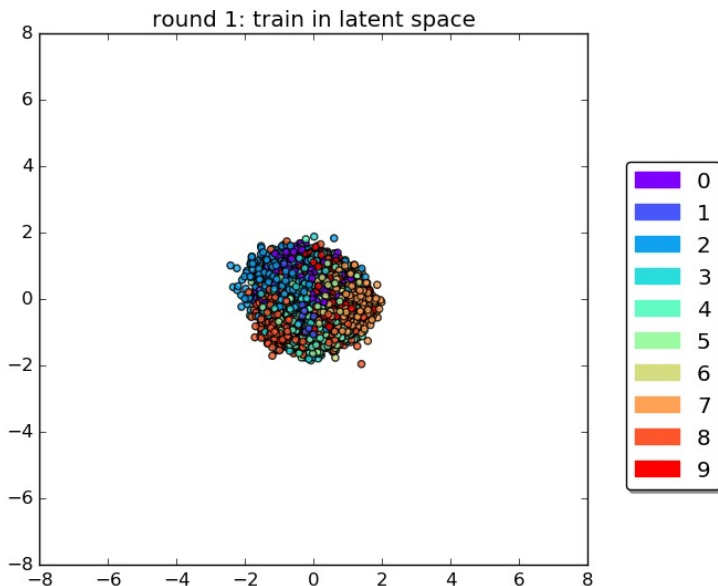
Example: 2D latent space on MNIST



“Introducing Variational Autoencoders (in Prose and Code)”

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

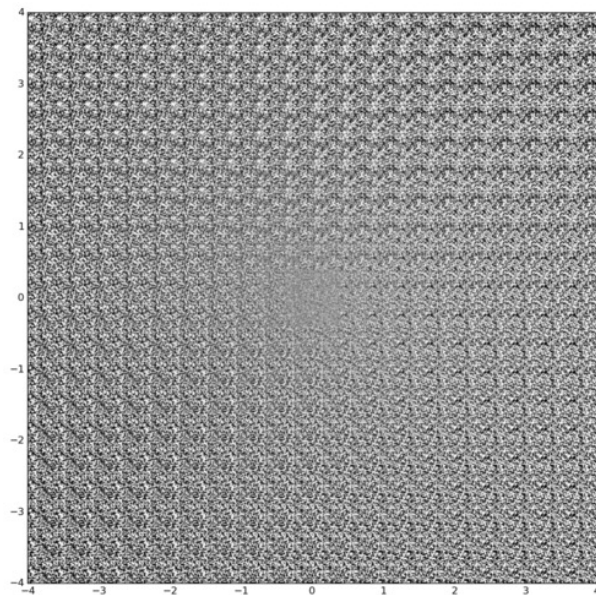
Example: 2D latent space on MNIST



“Introducing Variational Autoencoders (in Prose and Code)”

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

Example: 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>

Overview

- Introduction
 - What can Generative Models do?
 - What are Generative Models?
 - Generative models and representation learning
- Generative Models: Concepts and Recent Frontiers
 - Variational Autoencoders (VAE)
 - Diffusion models
- Embedded Ethics

Diffusion Models... in a nutshell

Forward process

- add noise to data

Reverse process

- learn to denoise

Training objective

- Denoising reconstruction loss

Diffusion Models... in a nutshell

Forward process

- add noise to data

Reverse process

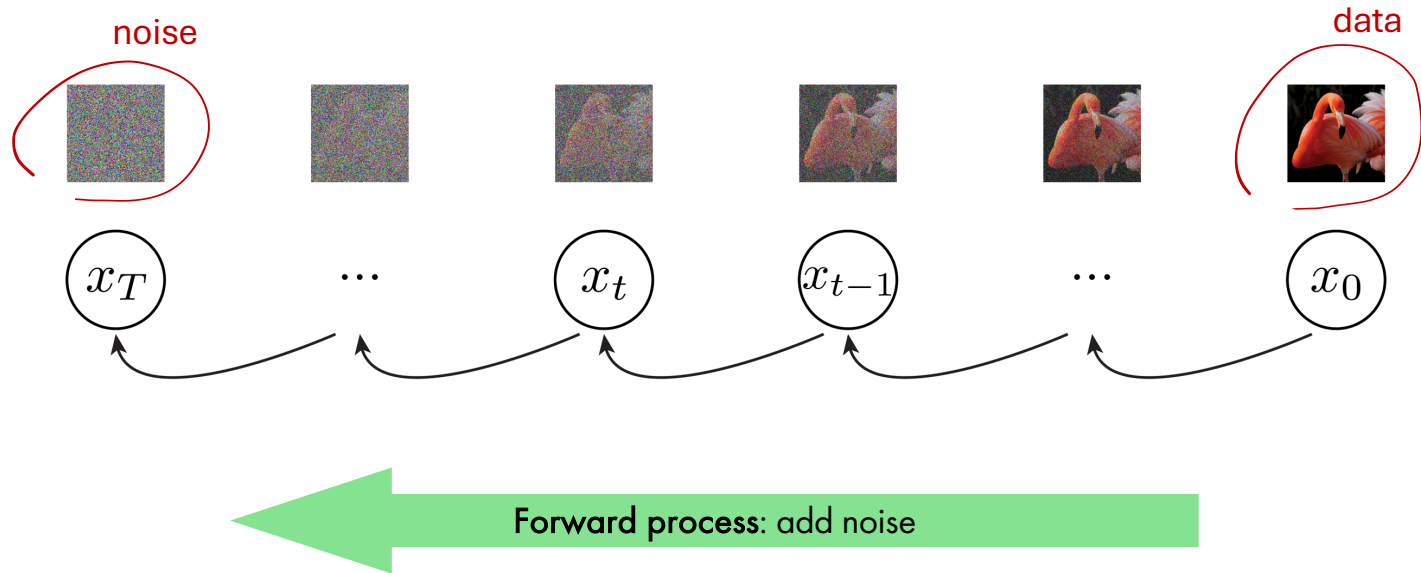
- learn to denoise

you'll derive these equations in your homework!

Training objective

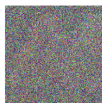
- Denoising reconstruction loss

Diffusion Models... in a nutshell



Diffusion Models... in a nutshell

noise



data

x_T

...

x_t

x_{t-1}

...

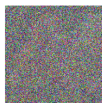
x_0

x_T is just standard normal distribution!

Forward process: add noise

Diffusion Models... in a nutshell

noise



data

x_T

...

x_t

x_{t-1}

...

x_0

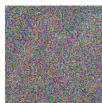
x_T is just standard normal distribution!

x_0 is also standard normal with zero variance and mean equal to the image pixel values!

Forward process: add noise

Diffusion Models... in a nutshell

noise



data



x_T

...

x_t

x_{t-1}

...

x_0

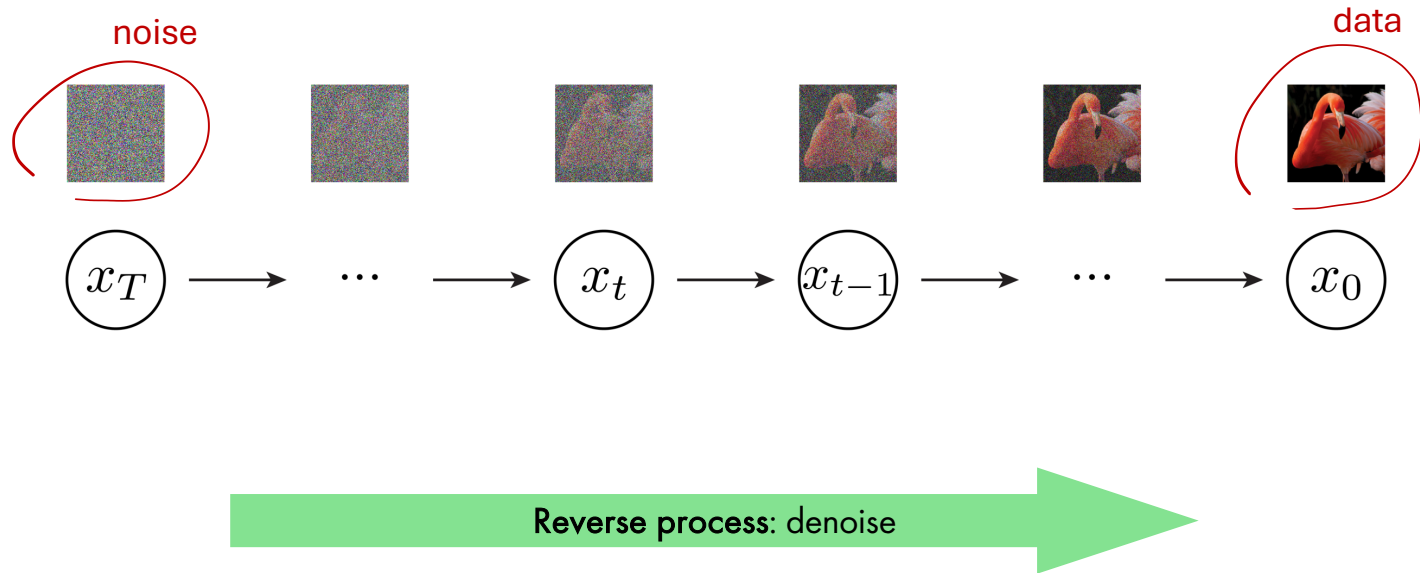
x_T is just standard normal distribution!

x_0 is also standard normal with zero variance and mean equal to the image pixel values!

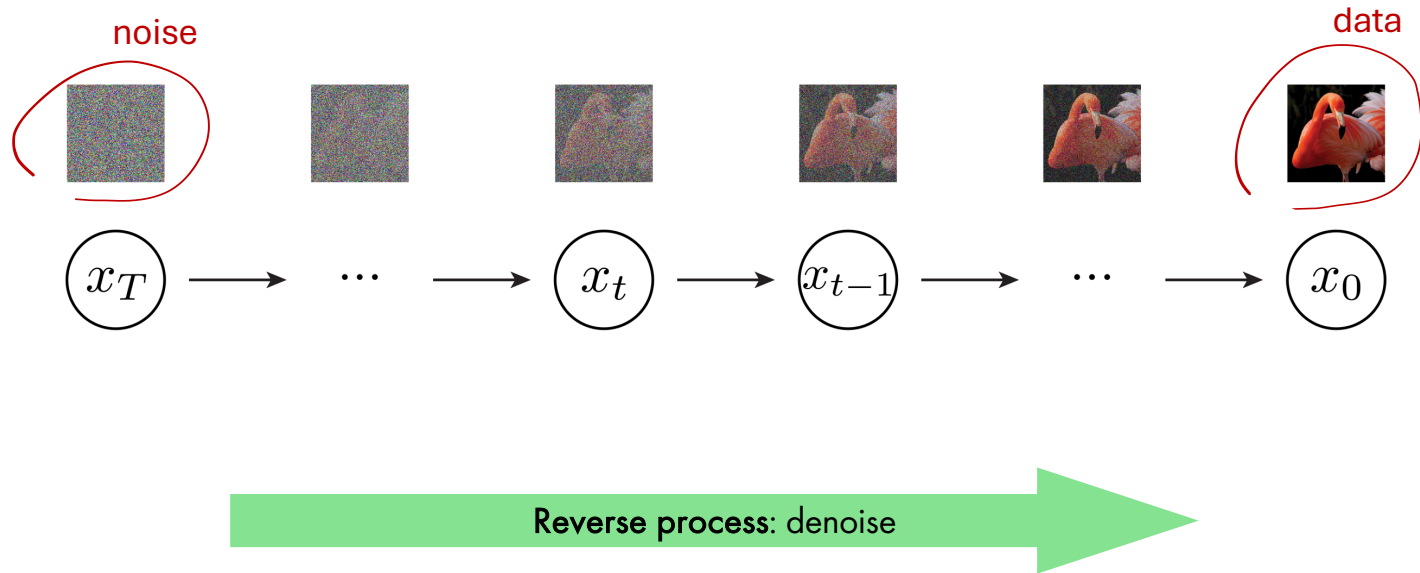
Forward process: add noise

so we can just add noise to the input image to find x_t for any value of t !

Diffusion Models... in a nutshell

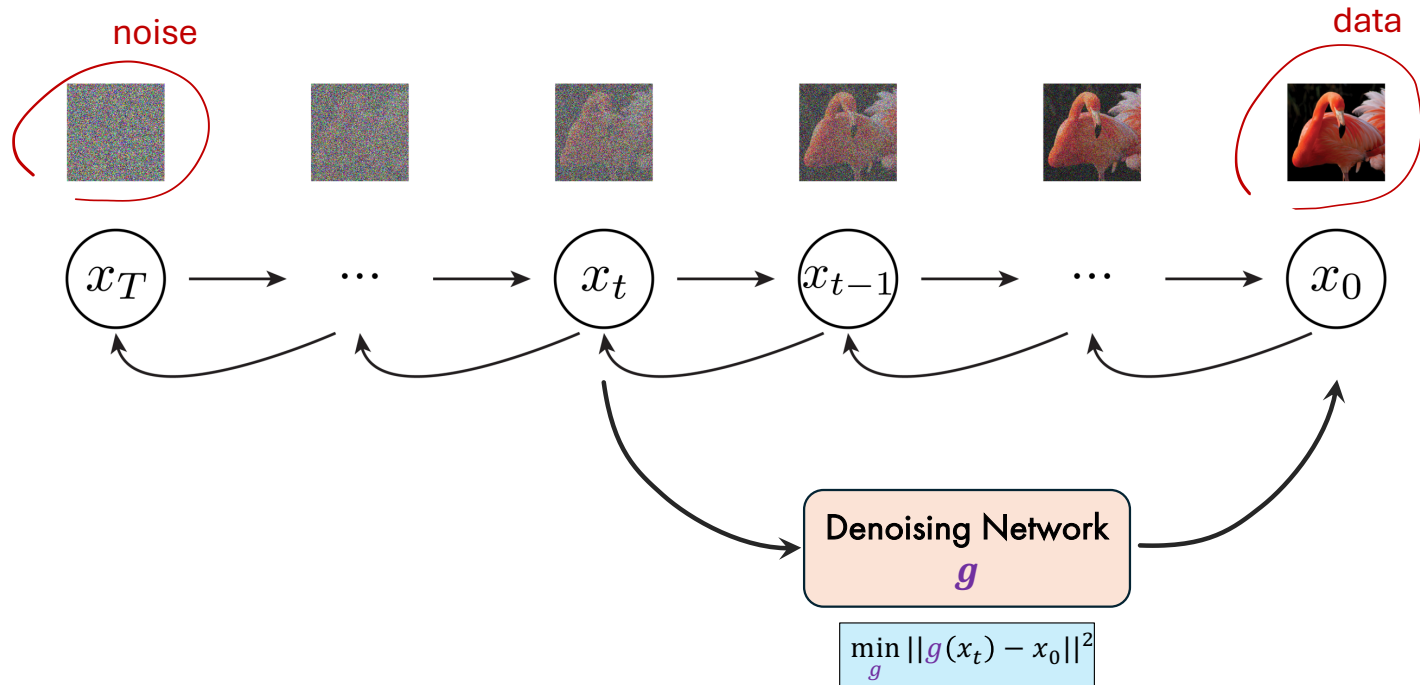


Diffusion Models... in a nutshell



reversing is a little bit harder — but we can show that if you know x_t and x_0 you can find an expression for x_{t-1} (i.e., the "slightly denoised image") analytically

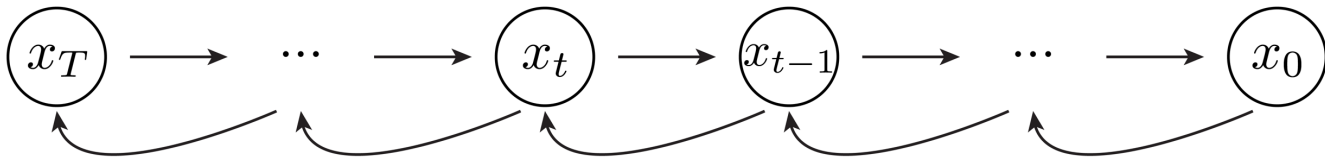
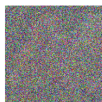
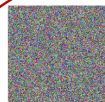
Diffusion Models... in a nutshell



Diffusion Models... in a nutshell

noise

data



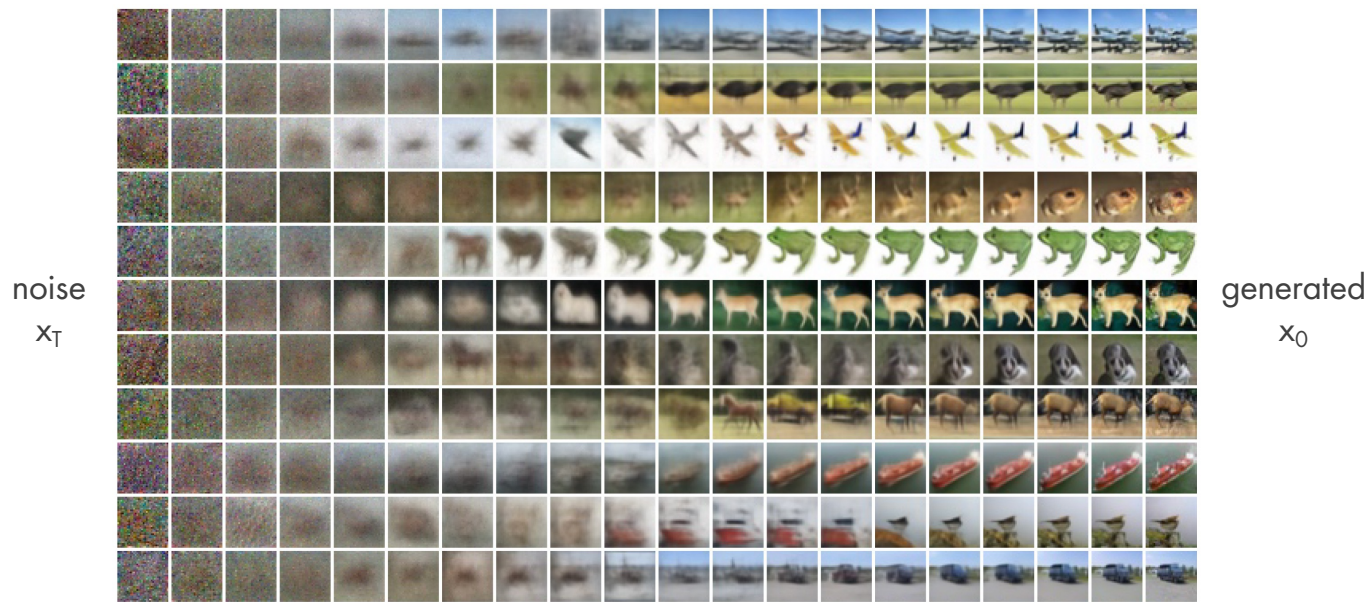
in practice the denoising network predicts the noise that was added — not the image
you'll show this in your homework!

Denoising Network

g

$$\min_g ||g(x_t) - x_0||^2$$

Example: Unconditional Generation on CIFAR-10



Overview

- Introduction
 - What can Generative Models do?
 - What are Generative Models?
 - Generative models and representation learning
- Generative Models: Concepts and Recent Frontiers
 - Variational Autoencoders (VAE)
 - Diffusion models
- Embedded Ethics

Warm-up activity



<https://www.cs.toronto.edu/~lindell/teaching/420/slides/celebahq.html>



Samples from CelebA-HQ Dataset

Warm-up activity



<https://www.cs.toronto.edu/~lindell/teaching/420/slides/celebahq.html>

Small group discussions

- What groups are not represented in this dataset?

Warm-up activity



<https://www.cs.toronto.edu/~lindell/teaching/420/slides/celebahq.html>

Small group discussions

- What groups are not represented in this dataset?
- What is this dataset not useful for?

Warm-up activity



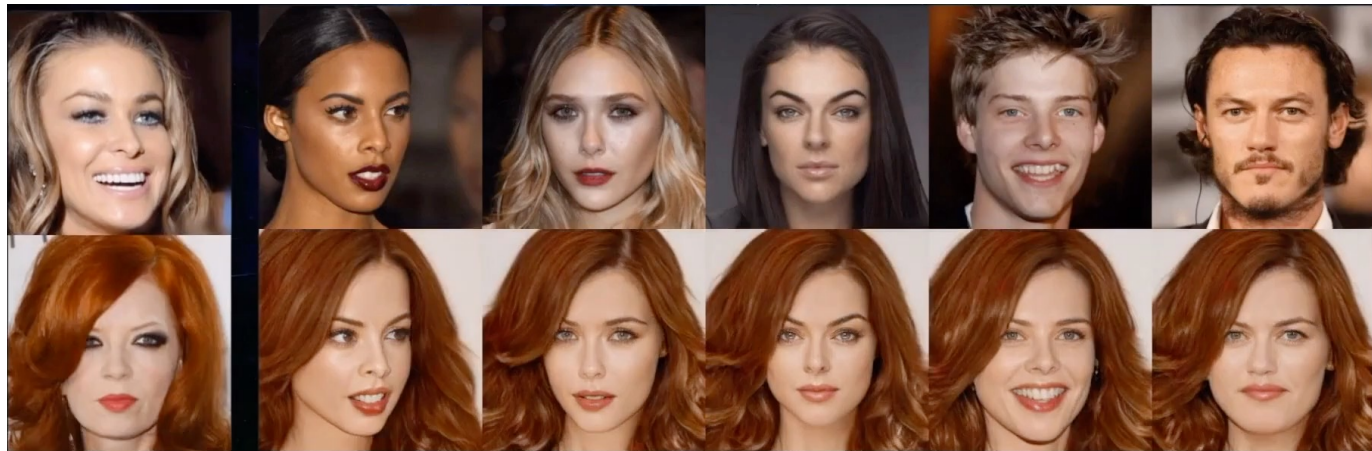
<https://www.cs.toronto.edu/~lindell/teaching/420/slides/celebahq.html>

Small group discussions

- What groups are not represented in this dataset?
- What is this dataset not useful for?
- What is this dataset useful for?

Image stylization with CelebA-HQ

source images



reference

stylized images

Projecting images onto CelebA-HQ models



[Esser et al. '20]

image generation in 2024

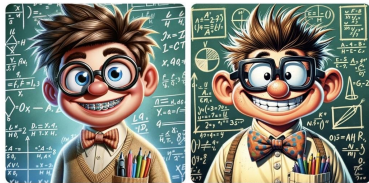


You

Make me an image of a cartoon nerd
without eyeglasses



DALL·E



Here are two images of a cartoon nerd
without eyeglasses. Let me know if
there's anything else you need!

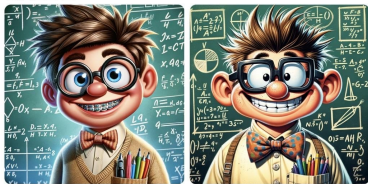


You

Make me an image of a cartoon nerd without eyeglasses



DALL-E



Here are two images of a cartoon nerd without eyeglasses. Let me know if there's anything else you need!

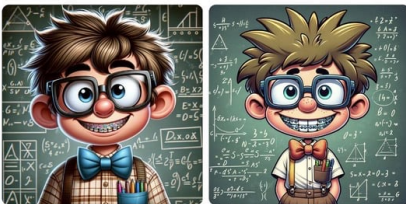


You

They both have eyeglasses on them , can you do them without eyeglasses



DALL-E



Here are two images of a cartoon nerd, created with the specific instruction to exclude eyeglasses. If you have any more requests or need further adjustments, feel free to let me know!

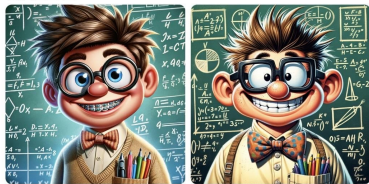


You

Make me an image of a cartoon nerd without eyeglasses



DALL·E



Here are two images of a cartoon nerd without eyeglasses. Let me know if there's anything else you need!

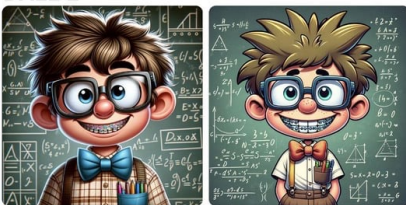


You

They both have eyeglasses on them , can you do them without eyeglasses



DALL·E



Here are two images of a cartoon nerd, created with the specific instruction to exclude eyeglasses. If you have any more requests or need further adjustments, feel free to let me know!



You

Ok, do you know what eyeglasses are ?



DALL·E

Yes, I understand what eyeglasses are. They are frames bearing lenses worn in front of the eyes, typically for vision correction or eye protection. I apologize for the oversight in the images provided. If you need an image without eyeglasses, I can certainly create one that meets your specifications.



You

Make me an image of only eyeglasses



DALL·E



image generation in 2025

make me an image of a cartoon nerd without eyeglasses



Here's your cartoon nerd without eyeglasses! Let me know if you'd like any tweaks.





Here's your cartoon nerd without eyeglasses! Let me know if you need any modifications.



no, it still has eyeglasses. please make the nerd without any kind of eye glasses. you do know what eyeglasses are, right?



Here you go—this nerd has no eyeglasses at all! Let me know if you'd like any adjustments.

Ethics of Bias & Discrimination

Mitigation Strategies

Collect Better Datasets?

- All datasets are finite attempts at sampling a distribution

Collect Better Datasets?

- All datasets are finite attempts at sampling a distribution
- Deep learning models used to obtain SOTA results on standard benchmarks are typically trained on hundreds of thousands to billions of training examples

Collect Better Datasets?

- All datasets are finite attempts at sampling a distribution
- Deep learning models used to obtain SOTA results on standard benchmarks are typically trained on hundreds of thousands to billions of training examples
- It is not always practical to collect representative examples of this size from various populations to include in the training data



coffee



ceramic



clipart



design



plastic



beer



tea



printed



travel



drawing



cup



creative



cute



🇸🇪 IKEA · In stock
IKEA - FÄRGKLAR Mug Matte ...



🇺🇸 Amazon.ca · In stock
20 OZ Large Coffee Mug, Smilat...



🇸🇪 IKEA · In stock
IKEA - DINERA Mug Dark gray...



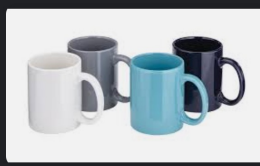
🇺🇸 Amazon.ca
New Ember Temperature ...



🇸🇪 IKEA
Mug, glossy beige, 13 oz (37 ...



🇺🇸 Apple · In stock
Ember Mug 2 Black 10 oz



🇨🇦 Canadian Tire
MASTER Chef 4pc Stoneware Mug Set ...



🇸🇪 IKEA · In stock
IKEA 365+ Mug



🇨🇦 YETI Canada · In stock
Yeti Rambler 414ml / 14oz Mug ...



🇨🇦 Rogue Wave Coffee · In stock
ORIGAMI - Barrel Aroma Mug Gr...



🇨🇦 Chapters Indigo · In stock
Friends Cappuccino Mug Centra...



🇨🇦 Main and Local · In stock
CBC Retro Logo Mug - Main and...



🇨🇦 Wayfair Canada · In stock
Coffee Mugs & Tea Cups - Wayf...



🇺🇸 Amazon.ca · In stock
Cricut Beveled Blank Mug, Cera...



🇨🇦 Hermès
Rocabar mug | Hermès Canada



🇺🇸 Williams Sonoma
Gold Monogram



🇨🇦 Vistaprint



🇺🇸 Apple



🇨🇦 The Toronto Star



🇨🇦 GameStop



🇨🇦 Canadian Tire



🇨🇦 Boutique RICARDO



🇨🇦 Nespresso · In stock

Gender Shades: Intersectional Accuracy Disparities in Commercial Gender Classification*

Joy Buolamwini

MIT Media Lab 75 Amherst St. Cambridge, MA 02139

JOYAB@MIT.EDU

Timnit Gebru

Microsoft Research 641 Avenue of the Americas, New York, NY 10011

TIMNIT.GEBRU@MICROSOFT.COM

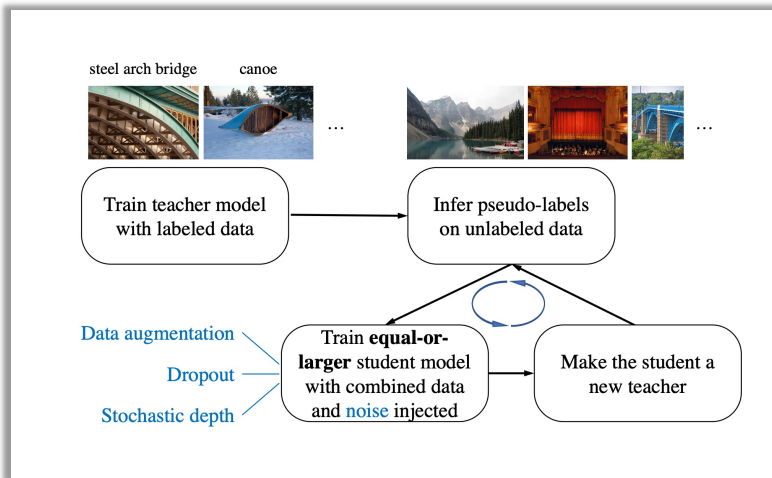
Editors: Sorelle A. Friedler and Christo Wilson

4.1. Key Findings on Evaluated Classifiers

- All classifiers perform better on male faces than female faces (8.1% – 20.6% difference in error rate)
- All classifiers perform better on lighter faces than darker faces (11.8% – 19.2% difference in error rate)
- All classifiers perform worst on darker female faces (20.8% – 34.7% error rate)
- Microsoft and IBM classifiers perform best on lighter male faces (error rates of 0.0% and 0.3% respectively)
- Face++ classifiers perform best on darker male faces (0.7% error rate)
- The maximum difference in error rate between the best and worst classified groups is 34.4%

Other mitigation strategies

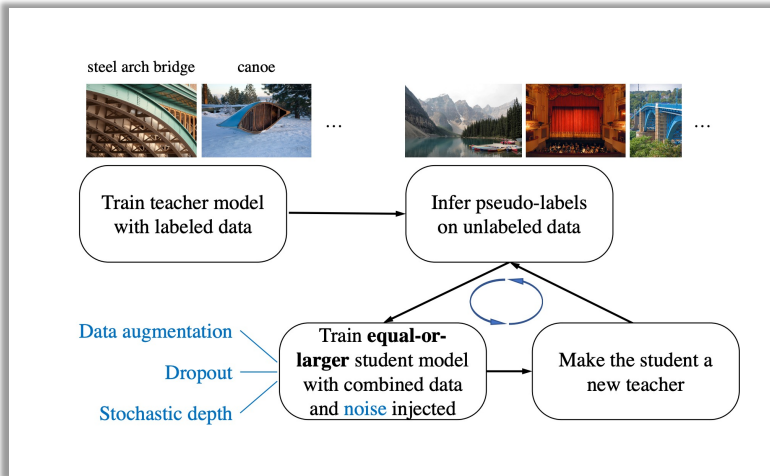
- Train on balanced dataset, finetune on biased dataset to reduce chances of overfitting



ResNet-50 Billion-scale [93]	26M	3.5B images labeled with tags	81.2%	96.0%
ResNeXt-101 Billion-scale [93]	193M		84.8%	-
ResNeXt-101 WSL [55]	829M		85.4%	97.6%
FixRes ResNeXt-101 WSL [86]	829M		86.4%	98.0%
Big Transfer (BiT-L) [43] [†]	928M	300M weakly labeled images from JFT	87.5%	98.5%
Noisy Student Training (EfficientNet-L2)	480M	300M unlabeled images from JFT	88.4%	98.7%

Other mitigation strategies

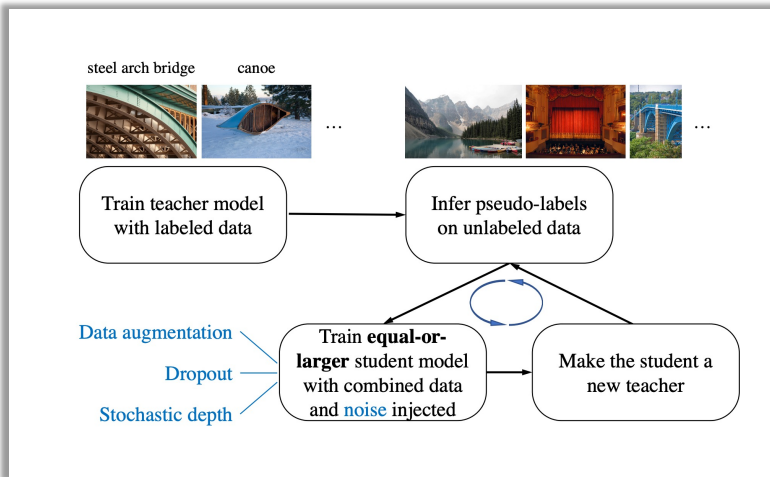
- Train on balanced dataset, finetune on biased dataset to reduce chances of overfitting
- Leverage labeled and unlabeled data to reduce sampling bias, increase size of dataset



ResNet-50 Billion-scale [93]	26M	3.5B images labeled with tags	81.2%	96.0%
ResNeXt-101 Billion-scale [93]	193M		84.8%	-
ResNeXt-101 WSL [55]	829M		85.4%	97.6%
FixRes ResNeXt-101 WSL [86]	829M		86.4%	98.0%
Big Transfer (BiT-L) [43] [†]	928M	300M weakly labeled images from JFT	87.5%	98.5%
Noisy Student Training (EfficientNet-L2)	480M	300M unlabeled images from JFT	88.4%	98.7%

Other mitigation strategies

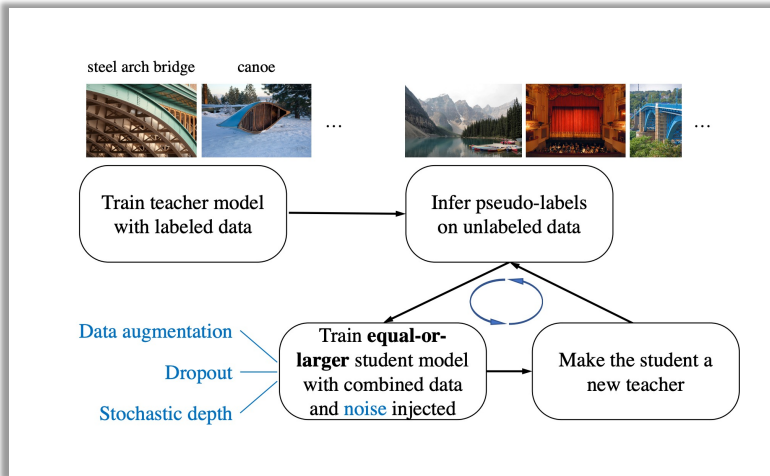
- Train on balanced dataset, finetune on biased dataset to reduce chances of overfitting
- Leverage labeled and unlabeled data to reduce sampling bias, increase size of dataset
- Data augmentation, re-balancing



ResNet-50 Billion-scale [93]	26M	3.5B images labeled with tags	81.2%	96.0%
ResNeXt-101 Billion-scale [93]	193M		84.8%	-
ResNeXt-101 WSL [55]	829M		85.4%	97.6%
FixRes ResNeXt-101 WSL [86]	829M		86.4%	98.0%
Big Transfer (BiT-L) [43] [†]	928M	300M weakly labeled images from JFT	87.5%	98.5%
Noisy Student Training (EfficientNet-L2)	480M	300M unlabeled images from JFT	88.4%	98.7%

Other mitigation strategies

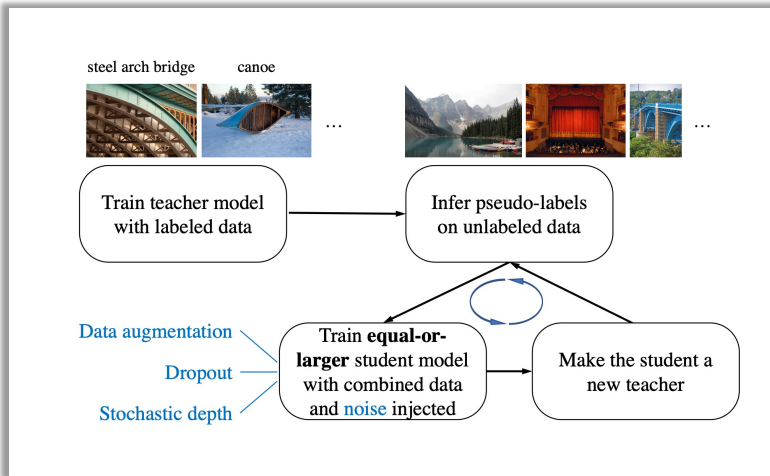
- Train on balanced dataset, finetune on biased dataset to reduce chances of overfitting
- Leverage labeled and unlabeled data to reduce sampling bias, increase size of dataset
- Data augmentation, re-balancing
- Evaluate on multiple datasets, balanced test sets



ResNet-50 Billion-scale [93]	26M	3.5B images labeled with tags	81.2%	96.0%
ResNeXt-101 Billion-scale [93]	193M		84.8%	-
ResNeXt-101 WSL [55]	829M		85.4%	97.6%
FixRes ResNeXt-101 WSL [86]	829M		86.4%	98.0%
Big Transfer (BiT-L) [43] [†]	928M	300M weakly labeled images from JFT	87.5%	98.5%
Noisy Student Training (EfficientNet-L2)	480M	300M unlabeled images from JFT	88.4%	98.7%

Other mitigation strategies

- Train on balanced dataset, finetune on biased dataset to reduce chances of overfitting
- Leverage labeled and unlabeled data to reduce sampling bias, increase size of dataset
- Data augmentation, re-balancing
- Evaluate on multiple datasets, balanced test sets
- ... (potentially many others!)



ResNet-50 Billion-scale [93]	26M	3.5B images labeled with tags	81.2%	96.0%
ResNeXt-101 Billion-scale [93]	193M		84.8%	-
ResNeXt-101 WSL [55]	829M		85.4%	97.6%
FixRes ResNeXt-101 WSL [86]	829M		86.4%	98.0%
Big Transfer (BiT-L) [43] [†]	928M	300M weakly labeled images from JFT	87.5%	98.5%
Noisy Student Training (EfficientNet-L2)	480M	300M unlabeled images from JFT	88.4%	98.7%

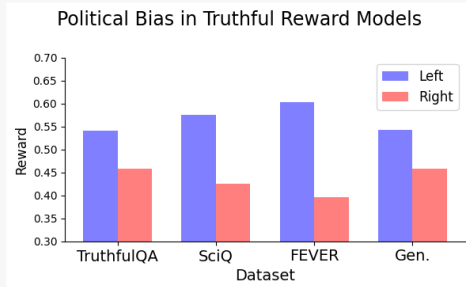
Takeaways

- dataset bias & mitigations still active research area!
- not only algorithmic performance, there are ethical implications
- depends both on data, but model architecture and even hyperparameter choices can all play a role

Study: Some language reward models exhibit political bias

Research from the MIT Center for Constructive Communication finds this effect occurs even when reward models are trained on factual data.

Ellen Hoffman | Center for Constructive Communication
December 10, 2024



Next time: Digital Photography