

11-695: AI Engineering

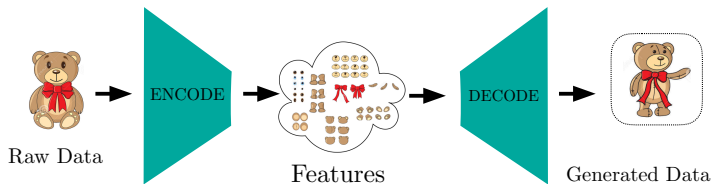
GAN

Spring 2019

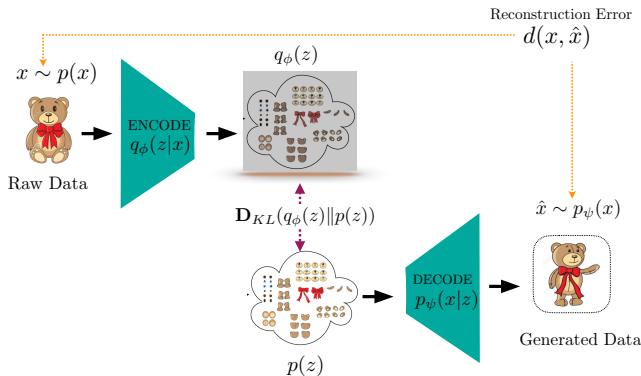
① AutoEncoder: Review and Motivation

② Minimax Game

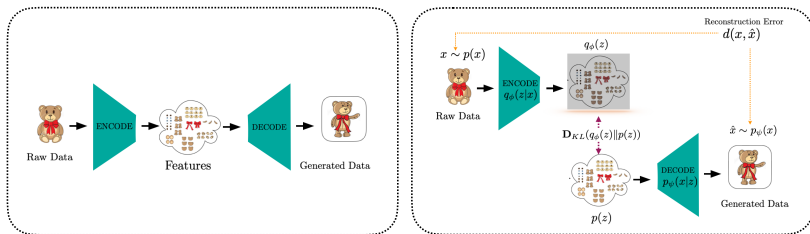
③ Generative Adversarial Networks



- We only have reconstruction loss
- Read more: Denoising AutoEncoder

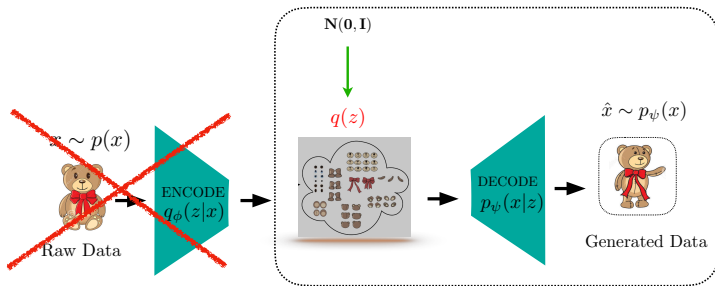


- Still AutoEncoder, with KL term for variational vs. true distribution $\rightarrow$ Variational AutoEncoder
- We have ELBO loss which is: Likelihood + KL distance

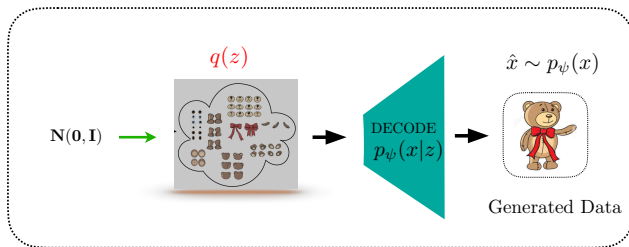


- Both approaches depend on how well you extract data, hence Likelihood-based approach
- With help from a noise (e.g. $\mathbf{N}(\mathbf{0}, \mathbf{I})$), data can be better generated

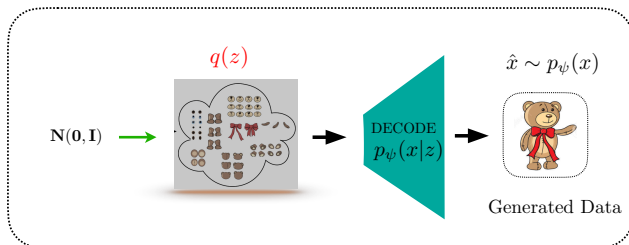
What about no likelihood?



- What if we get rid of Encode part from data
- *i.e.* we are left with generation only noise space
- Recall: We need to tune the noise. How?
- Recall: Use KL distance from a variational distribution to the true distribution

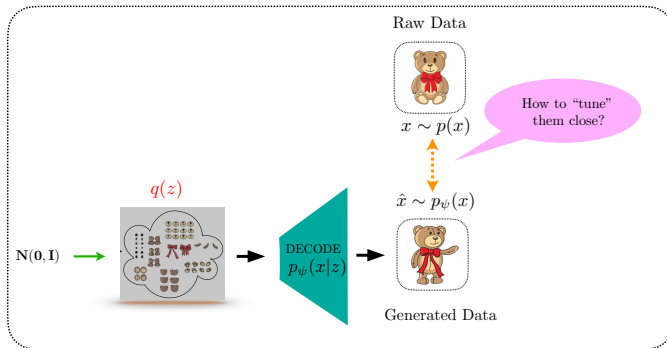


- $ELBO = \log p(x) - \mathbf{D}_{KL}(q(z) \| p(z|x))$
- But we don't have $p(z|x)$ and so no base distribution to do a variational approximation on
- But: we need to make use of data ($x \sim p(x)$) to tune Decoder



- $ELBO = \log p(x) - \mathbf{D}_{KL}(q(z) \| p(z|x))$
- But we don't have $p(z|x)$ and so no base distribution to do a variational approximation on
- And we generate from a noise (random) space

Replace Encoder with a New Component

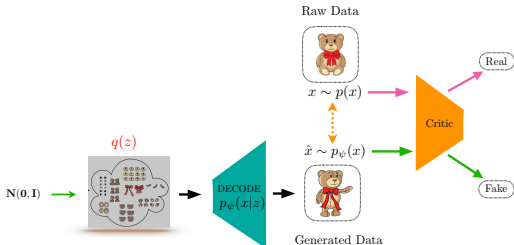


- Why don't we try to have a critic to evaluate how well the generator is?
- (hence the "Adversarial")

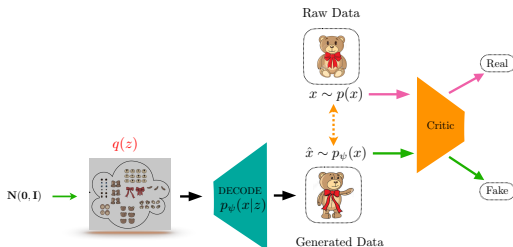
① AutoEncoder: Review and Motivation

② Minimax Game

③ Generative Adversarial Networks

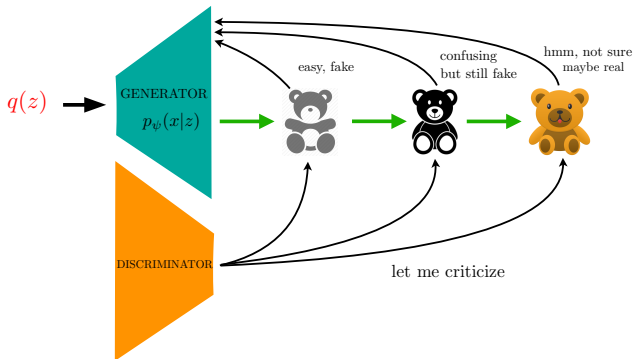


- It needs to make use of raw data, the only resource we have, and to compare with.
- Design: like encoder, it should be able to extract features from data (of course, normally with a NN)
 - But only for a much simpler task: binary classification
 - So it looks like we have shifted encoder from front to back



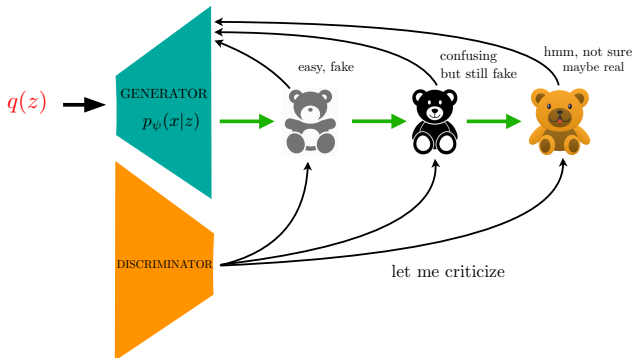
- Binary classification: isn't it too simple?
- Not quite, if we include a feedback loop, in that
 - Assume Critic is strict, and perfect, and constantly gives “feedback”
 - Generator uses feedback to get better overtime

Call them Generator and Discriminator

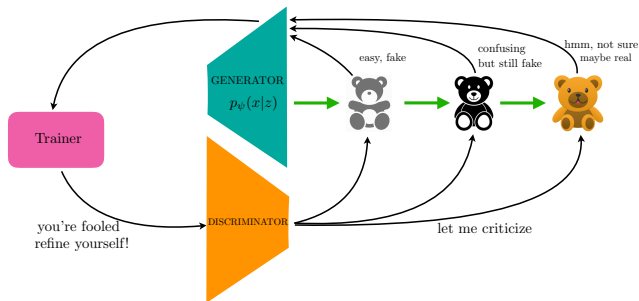


- Binary classification: isn't it too simple?
- Not quite, if we include a feedback loop, in that
 - Assume Critic is strict, and perfect, and constantly gives "feedback"
 - Generator uses feedback to get better overtime

We call the feedback process is a game



- Generator tries to fool the Discriminator
- Discriminator tries not to be fooled by Generator



- For sure we need! It's a NN and thus needs training
- As generator is better, discriminator has to be better too
- So both end up in a state where generated images are hard to classify

- Data x , Noise z , Generator G , Discriminator D
 - Real data: $x \sim p(x)$
 - Generated (fake) data: $\hat{x} \sim p(z) = G_\psi(z)$
 - Discriminator classify $D_\theta(a)$ with input a
- Discriminator to, at the same time, **maximize**:
 - to tune $D_\theta(x)$ close to 1,
 - and $D_\theta(G_\psi(z))$ close to 0
- Generator, conversely, wants to **minimize** $D_\theta(G_\psi(z))$ to 1

- Discriminator to **maximize**:

- to tune $D_\theta(x)$ close to 1: $\mathbb{E}_{x \sim p(x)} \log D_\theta(x)$
- and $D_\theta(G_\psi(z))$ close to 0: $\mathbb{E}_{z \sim p(z)} \log D_\theta(G_\psi(z))$

- Rewrite Discriminator:

$$\max_{\theta} [\mathbb{E}_{x \sim p(x)} \log D_\theta(x) + (1 - \mathbb{E}_{z \sim p(z)} \log D_\theta(G_\psi(z)))]$$

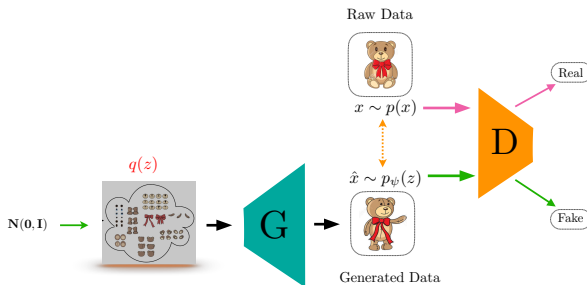
- Generator: tries to **minimize** that term:

$$\min_{\psi} \max_{\theta} [\mathbb{E}_{x \sim p(x)} \log D_\theta(x) + (1 - \mathbb{E}_{z \sim p(z)} \log D_\theta(G_\psi(z)))]$$

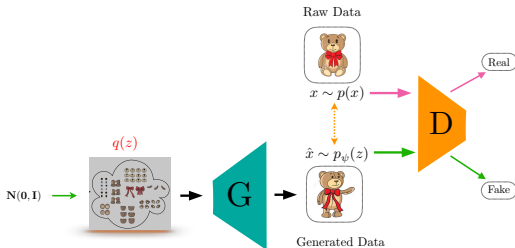
① AutoEncoder: Review and Motivation

② Minimax Game

③ Generative Adversarial Networks



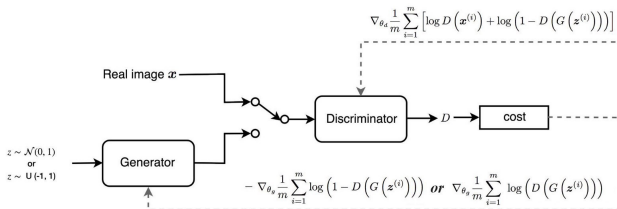
- A Generator: to fool a Discriminator
- A Discriminator: to play against Generator, hence “adversarial”



- Objective:

$$\min_{\psi} \max_{\theta} [\mathbb{E}_{x \sim p(x)} \log D_{\theta}(x) + (1 - \mathbb{E}_{z \sim p(z)} \log D_{\theta}(G_{\psi}(z)))]$$

- We alternate in that:
 - Make G constant, let D **maximize** the objective
 - Make D constant, let G **minimize** the objective
- Question: single backprop?



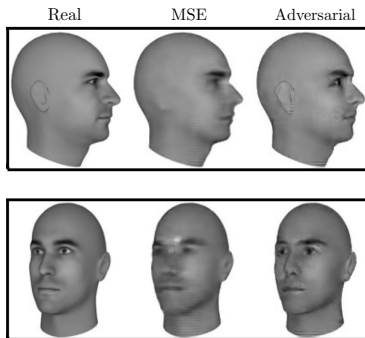
- Empirical expectation for gradient of D_θ :

$$\nabla_\theta = \frac{1}{m} \sum_{i=1}^m [\log D_\theta(x^{(i)}) + \log(1 - D_\theta(G_\psi(z^{(i)})))]$$

- Empirical expectation for gradient of G_ψ :

$$\nabla_\psi = \frac{1}{m} \sum_{i=1}^m [\log(1 - D_\theta(G_\psi(z^{(i)})))]$$

- Batching for D_θ : half real and half fake.



- MLE: blurry, seems to have the average effect
- GAN: sharp, seems to directly maximize the real-like generation (thanks to minimax theory)

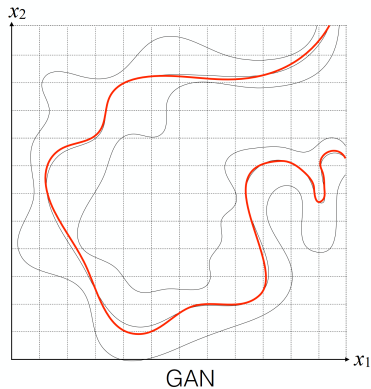
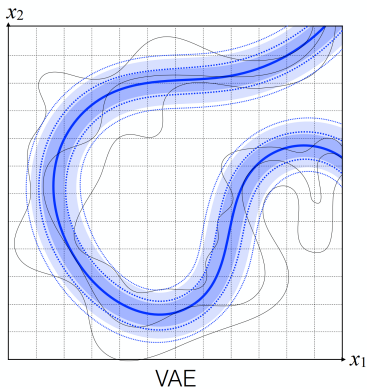




Figure 5: 1024×1024 images generated using the CELEBA-HQ dataset. See Appendix F for a larger set of results, and the accompanying video for latent space interpolations.