



· Likelihood and Sample Quality

1. high log-likelihood, poor samples

$$P_{\theta}(x) = 0.01 \cdot P_{\text{data}}(x) + 0.99 \cdot P_{\text{noise}}(x)$$

99% samples are noise

$$\begin{aligned} \log P_{\theta}(x) &= \log [0.01 P_{\text{data}}(x) + 0.99 P_{\text{noise}}(x)] \\ &\geq \log 0.01 P_{\text{data}}(x) \\ &= \log P_{\text{data}}(x) - \log 100 \end{aligned}$$

lower bound: $E_{x \sim P_{\theta}}[P_{\theta}(x)] \geq E_{x \sim P_{\theta}}[\log P_{\text{data}}(x)] - \log 100$

upper bound: $KL(P_{\text{data}} \| P_{\theta}) = E_{x \sim P_{\text{data}}}[\log P_{\text{data}}(x)] - E_{x \sim P_{\theta}}[\log P_{\theta}(x)] \geq 0$

$$E_{x \sim P_{\theta}}[P_{\theta}(x)] \leq E_{x \sim P_{\theta}}[\log P_{\text{data}}(x)]$$

$$\begin{array}{ccc} \text{---} E_{x \sim P_{\theta}}[\log P_{\text{data}}(x)] \text{---} & & \downarrow \\ & E_{x \sim P_{\theta}}[P_{\theta}(x)] & 100 \\ & & \uparrow \\ \text{---} E_{x \sim P_{\theta}}[\log P_{\text{data}}(x)] - 100 \text{---} & & \end{array}$$

as we increase the dimensionality of x , $\log P_{\text{data}}(x)$ gets bigger in absolute value
then in the sense of relative difference, $E_{x \sim P_{\theta}}[\log P_{\theta}(x)] \approx E_{x \sim P_{\text{data}}}[\log P_{\text{data}}(x)]$

high log-likelihood does not imply good sampling

2. good samples, low likelihoods $P_{\theta}(x_{\text{sample}})$
just memorize the training set

Two-Sample Test

given $S_1 = \{X \sim P\}$ $S_2 = \{X \sim Q\}$, a two-sample test consider the following hypothesis

null hypothesis: $H_0: P=Q$

alternate hypothesis: $H_1: P \neq Q$

Test statistic T compares S_1 and S_2 (eg. difference in mean or variance)

if T is less than a threshold α , accept H_0

Test statistic T is likelihood-free since it does not involve $P(X)$ or $Q(X)$

Assume direct access to $S_1 = D = \{X \sim P_{data}\}$

have a model distribution P_0 which is easy to sample from, let $S_2 = \{X \sim P_0\}$

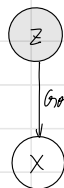
train the generative model to minimize a two-sample test objective between S_1 and S_2

Test statistic: Learn a statistic that maximize a suitable notion of distance between S_1 and S_2

Generator Adversarial Network

generator

1. directed latent variable model with a deterministic mapping from z to x given by G_0
2. minimizes a two-sample test objective (in support of null hypothesis $P_{data} = P_0$)



discriminator:

1. any function (eg. neural network) which tries to distinguish "real" samples from the dataset and "fake" samples from generator
2. maximize the two-sample test objective (in support of the alternate hypothesis $P_{data} \neq P_0$)



discrimination: $\max_D V(G, D)$

$$V(G, D) = \underbrace{E_{x \sim P_0} [\log D(x)]}_{\text{want a large log probability for samples from data distribution}} + \underbrace{E_{x \sim P_G} [\log (1 - D(x))]}_{\text{want a low log probability for samples from generator distribution}}$$

want a large log probability
for samples from data distribution

want a low log probability
for samples from generator distribution

optimal discriminator: $D_0^*(x) = \frac{P_{data}(x)}{P_{data}(x) + P_0(x)}$

$$V(G, D) = \int_{\mathcal{X}} P_{data}(x) \log D(x) dx + \int_{\mathcal{X}} P_0(x) \log (1 - D(x)) dx$$

$$\frac{dV}{dD_0(x)} = P_{data}(x) \cdot \frac{1}{D(x)} - P_0(x) \cdot \frac{1}{1 - D(x)} = 0$$

$$D_0^*(x) = \frac{P_{data}(x)}{P_{data}(x) + P_0(x)}$$

if $P_{data}(x) \approx P_0(x)$,

then the discriminator cannot tell whether x
is from P_{data} or P_0

generator: $\min_G V(G, D)$
 $\max_D \mathbb{E}_{x \sim P_G} [\log(1 - D(x))]$

for the optimal discriminator $D_G^*(\cdot)$, we have

$$\begin{aligned} V(G, D_G^*) &= \mathbb{E}_{x \sim P_{data}} [\log D_G^*(x)] + \mathbb{E}_{x \sim P_G} [\log(1 - D_G^*(x))] \\ &= \mathbb{E}_{x \sim P_{data}} \left[\log \frac{P_{data}(x)}{P_{data}(x) + P_G(x)} \right] + \mathbb{E}_{x \sim P_G} \left[\log \frac{P_G(x)}{P_{data}(x) + P_G(x)} \right] \\ &= \mathbb{E}_{x \sim P_{data}} \left[\log \frac{P_{data}(x)}{\frac{P_{data}(x) + P_G(x)}{2}} \right] + \mathbb{E}_{x \sim P_G} \left[\log \frac{P_G(x)}{\frac{P_{data}(x) + P_G(x)}{2}} \right] - \log 4 \\ &= \underbrace{KL \left[P_{data} \parallel \frac{P_{data} + P_G}{2} \right] + KL \left[P_G \parallel \frac{P_{data} + P_G}{2} \right]}_{2 \times \text{Jensen-Shannon Divergence}} - \log 4 \\ &= 2 \times D_{JSD}(P_{data} \parallel P_G) - \log 4 \end{aligned}$$

the optimal generator approximate $P_G = P_{data}$

for the optimal generator $G^*(\cdot)$ and $D_G^*(\cdot)$, we have $V(G^*, D_G^*) = -\log 4$

Jensen-Shannon Divergence

$$D_{JSD}(P \parallel Q) = \frac{1}{2} \left[KL(P \parallel \frac{P+Q}{2}) + KL(Q \parallel \frac{P+Q}{2}) \right]$$

also called symmetric KL divergence

$$D_{JSD}(P \parallel Q) \geq 0 \quad \Rightarrow \quad \text{iff } P=Q$$

$$D_{JSD}(P \parallel Q) = D_{JSD}(Q \parallel P)$$

D_{JSD} is not a measure of distance

since it does not satisfy the triangle inequality, but JD_{JSD} does

GAN Training Algorithm

Sample $x^{(1)} \dots x^{(m)}$ from D

Sample $z^{(1)} \dots z^{(m)}$ from noise P_z

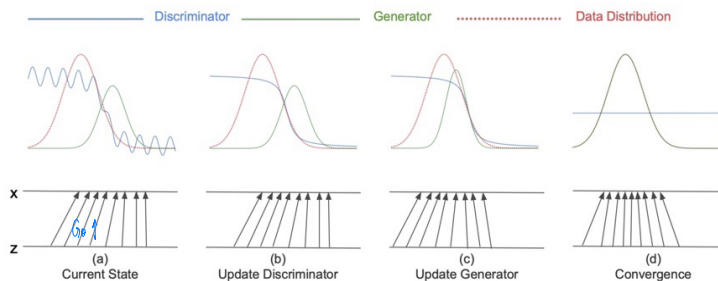
Update the generator parameters θ by stochastic gradient *descent*

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

Update the discriminator parameters ϕ by stochastic gradient *ascent*

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

$$V(G, D) = \mathbb{E}_{x \sim p_{data}} [\log D(x)] + \mathbb{E}_{z \sim P_z} [\log 1 - D(G(z))]$$



Optimization Challenges

Unrealistic assumption: we specify a function and only make update in parameter space

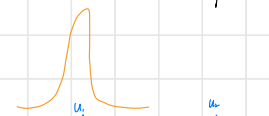
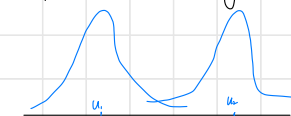
if the generator updates are made in function space, and discriminator is optimal at every step, then the generator guaranteed to converge to the data distribution *unrealistic*

no robust stopping criteria in practice (unlike MLE)

Mode Collapse

GANs are notorious for suffering from mode collapse

the phenomena where the generator collapses to one or several examples



Fixing mode collapse are mostly empirically driven:

alternate architectures, adding regularization, injecting small noise perturbation

<https://github.com/soumith/ganhacks>