

GAN: Generative Adversarial Network

$$\{x_i\}_{i=1}^M, p(z) \rightarrow x = g(z)$$

$\Rightarrow$ Learning $g: z \rightarrow x$

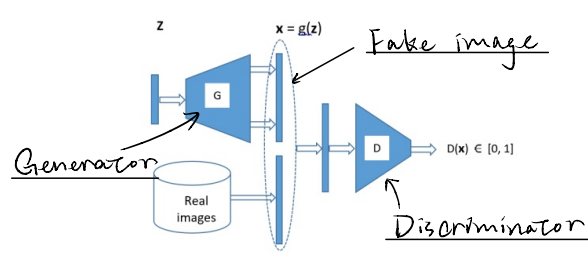
Training G, D Alternatively.

Training D $\Rightarrow$ tell real & fake.

$\{D(x) = 1 \text{ or } \rightarrow 1 : \text{real.}$

$\{D(x) = 0 \text{ or } \rightarrow 0 : \text{fake}$

Training G $\Rightarrow$ fool D



θ_g : parameter for G. θ_d : parameter for D

1 Improve the θ_d so that the discriminator becomes better at distinguishing between fake and real examples: Run the following k times:

- Sample m real examples $x_1, \dots, x_m$ from training data.
- Generate m fake examples $g(z_1), \dots, g(z_m)$ using the current generator g , where $z_i \sim p(z)$.
- Improve θ_d so that the discriminator can better distinguish between those fake examples and those real examples.

2 Improve generator θ_g so as to generate examples the discriminator would find hard to tell fake or real:

- Generate m new fake examples $g(z_1), \dots, g(z_m)$ using the current generator g . (Those are examples that the discriminator can tell as fake with high confidence because of the training in step 1.)
- Improve θ_g to generate examples that would be more like real images than those fake images.

Training D:

$$\begin{aligned} &\text{real: } \{x_i\}_{i=1}^M \\ &\text{fake: } \{g(z_i)\}_{i=1}^M \end{aligned} \Rightarrow \begin{aligned} &\text{Loss (Cross Entropy)} \\ &-\log(\gamma = 1 | x_i) = -\log D(x_i) \\ &-\log(\gamma = 0 | g(z_i)) = -\log D(g(z_i)) \end{aligned}$$

$$\Rightarrow D \text{ minimize: } -\frac{1}{m} \sum_{i=1}^m \log D(x_i) - \frac{1}{m} \sum_{i=1}^m \log D(1 - g(z_i))$$

$$\Leftrightarrow \text{maximize: } \frac{1}{m} \sum_{i=1}^m \log D(x_i) + \frac{1}{m} \sum_{i=1}^m \log D(1 - g(z_i))$$

$$V(\theta_g, \theta_d) = \frac{1}{m} \sum_{i=1}^m \log D_{\theta_d}(x_i) + \frac{1}{m} \sum_{i=1}^m \log D(1 - g_{\theta_g}(z_i))$$

$$\text{Obj} \rightarrow D: \max_{\theta_d} V(\theta_g, \theta_d) = V(\theta_g)$$

(Minimax)

$$G: \min_{\theta_g} \max_{\theta_d} V(\theta_g, \theta_d) = \min_{\theta_g} V(\theta_g)$$

Training G:

$$\begin{aligned} \min_{\theta_g} V(\theta_g) &= \min_{\theta_g} V(\theta_g, \theta_d^*) \\ &= \frac{1}{m} \sum_{i=1}^m \log D^*(x_i) + \frac{1}{m} \sum_{i=1}^m \log(1 - D(g(z_i))) \\ &= \min_{\theta_g} \frac{1}{m} \sum_{i=1}^m \log(1 - D(g(z_i))) \end{aligned}$$

Does not depend on θ_d

Algorithm 1 Minibatch stochastic gradient descent training of generative adversarial nets. The number of steps to apply to the discriminator, k , is a hyperparameter. We used $k = 1$, the least expensive option, in our experiments.

for number of training iterations do

for k steps do

- Sample minibatch of m noise samples $\{z^{(1)}, \dots, z^{(m)}\}$ from noise prior $p_g(z)$.
- Sample minibatch of m examples $\{x^{(1)}, \dots, x^{(m)}\}$ from data generating distribution $p_{\text{data}}(x)$.
- Update the discriminator by ascending its stochastic gradient:

$$\nabla_{\theta_d} \frac{1}{m} \sum_{i=1}^m \left[\log D(x^{(i)}) + \log(1 - D(g(z^{(i)}))) \right].$$

end for

- Sample minibatch of m noise samples $\{z^{(1)}, \dots, z^{(m)}\}$ from noise prior $p_g(z)$.
- Update the generator by descending its stochastic gradient:

$$\nabla_{\theta_g} \frac{1}{m} \sum_{i=1}^m \log(1 - D(g(z^{(i)}))).$$

end for

- At each iteration, the discriminator is not trained to optimum. Instead, its parameters are improved only once by gradient ascent.

- Similarly, the parameters of the generators are also improved only once in each iteration by gradient descent.

FID

- FID compares how often the same high-level features are found within the real and generate.
- Lower FID score means that the generated images have better quality and diversity.
- FID is an improved of Inception Score.

Theory of GAN

D maximizes:

$$V(\theta_g, \theta_d) = \frac{1}{m} \sum_{i=1}^m \log D(x_i) + \frac{1}{m} \sum_{i=1}^m \log(1 - D(g(z_i)))$$

$$\begin{aligned} V(G, D) &= E_{x \sim p_r} [\log D(x)] + E_{x \sim p_g} [\log(1 - D(g(z_i)))] \\ &= \int [Pr(x) \log D(x) + P_g(x) \log(1 - D(x))] dx. \end{aligned}$$

when V is optimized (fixed G)

$$\begin{aligned} D^*(x) &= \frac{Pr(x)}{Pr(x) + P_g(x)} \\ V(G, D^*) &= \int Pr(x) \cdot \log \frac{Pr(x)}{Pr(x) + P_g(x)} + P_g(x) \log \frac{Pr(x)}{Pr(x) + P_g(x)} dx \\ &= \int Pr(x) \log \frac{Pr(x)}{Pr(x) + P_g(x)} dx + \int P_g(x) \log \frac{P_g(x)}{Pr(x) + P_g(x)} dx - 2 \log 2 \\ &= 2 \times \frac{1}{2} \times (D_{KL}(Pr || Pa) + D_{KL}(Pg || Pa)) - 2 \log 2. \end{aligned}$$

$$= 2 \boxed{D_{JS}(Pr || Pg)} - 2 \log 2$$

Jensen-Shannon Divergence

GAN: $\begin{cases} \text{Generator: minimize JS Divergence} \\ \text{Discriminator: approximate JS Divergence} \end{cases}$

Loss function for G:

$$L = \frac{1}{m} \sum_{i=1}^m \log(1 - D(g(z_i)))$$

Gradient $\nabla_{\theta_g} E_{x \sim p_g} [\log(1 - D(x))]$

$$= \int \log(1 - D(x)) \nabla_{\theta_g} P_g(x) dx.$$

At Beginning: $D(x) \rightarrow 0, \log(1 - D(x)) \rightarrow 0.$

$\Rightarrow$ small gradient $\Rightarrow$ slow training.

Alternative Loss for G

$$L = \frac{1}{m} \sum_{i=1}^m -\log D(g(z_i))$$

$$\nabla_{\theta_g} L = \nabla_{\theta_g} E_{x \sim p_g(x)} [-\log D(x)]$$

$$= - \int \log D(x) \nabla_{\theta_g} P_g(x) dx.$$

$\Rightarrow$ gradient is large at beginning $\Rightarrow$ training is fast

$\Rightarrow$ when D is good ($D(x) \approx 0$ when $x \sim p_g(x)$), gradient $\uparrow$, unstable

- The original generator cost function leads to slow training, and the alternative cost function leads to unstable training.
- In practice, those difficulties are dealt with by not training the discriminator to optimum, not even close to optimum.

VAE:

- Training is stable.
- Does not suffer from mode dropping because it maximizes likelihood, ensuring coverage of the data distribution.
- Outputs are often blurrier due to the reconstruction loss, which averages over possible outputs, sacrificing fine details.

GAN:

- Generates sharper and more realistic outputs because the discriminator pushes the generator to produce highly detailed outputs.
- Prone to mode dropping, where the discriminator does not force the generator to capture the full diversity of the data distribution.
- Training is unstable and can suffer from issues like non-convergence or mode collapse.

VAE: Image Generation $\rightarrow$ regression

GAN: Image Generation $\rightarrow$ classification