


f -divergence

given 2 probability distribution p and q , the f -divergence is

$$D_f(p||q) = \mathbb{E}_{x \sim q} \left[f\left(\frac{p(x)}{q(x)}\right) \right]$$

where f is convex, lower-semicontinuous and $f(1) = 0$

A function is lower-semicontinuous at x_0 if

$$\forall y < f(x_0), \exists \epsilon \text{ s.t. } f(x) > y \quad \forall x \in [x_0 - \epsilon, x_0 + \epsilon]$$

$f([x_0 - \epsilon, x_0 + \epsilon])$ should be close to or larger than $f(x_0)$

eg. $f(x) = x \cdot \log x$, the f -divergence is KL divergence

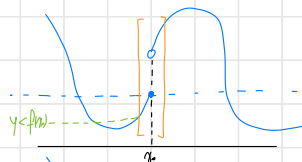
$$\mathbb{E}_{x \sim q} \left[\frac{p(x)}{q(x)} \cdot \log \frac{p(x)}{q(x)} \right] = \mathbb{E}_{x \sim p} \left[\log \frac{p(x)}{q(x)} \right]$$

for any function $f(\cdot)$, its conjugate is

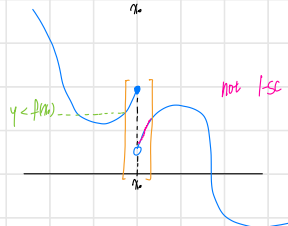
$$f^*(y) = \sup_{x \in \text{dom } f} (yx - f(x)) \quad f \text{ is convex}$$

$$f^{**} = f^{\text{conv}} \quad (\text{epi}(f^{**}) = \text{Convex hull}(\text{epi}(f)))$$

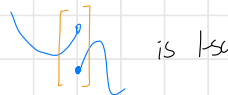
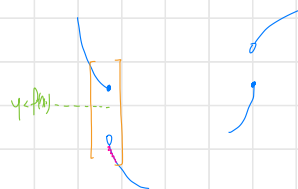
if f is convex, then $f^{**} = f$



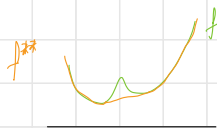
lower-semicontinuous at x_0



not lsc



is lsc



• f-GAN: Variational Divergence Minimization

Can obtain a lower bound to any f-divergence via its conjugate

$$\begin{aligned} D_f(P, Q) &= E_{x \sim q} \left[f\left(\frac{p(x)}{q(x)}\right) \right] \\ &= E_{x \sim q} \left[f^{**}\left(\frac{p(x)}{q(x)}\right) \right] \\ &= E_{x \sim q} \left[\sup_{y \in \text{dom} f^*} \left[y \cdot \frac{p(x)}{q(x)} - f^*(y) \right] \right] \\ &= E_{x \sim q} \left[T(x) \cdot \frac{p(x)}{q(x)} - f^*(T(x)) \right] \\ &= \int_{\mathcal{X}} T(x) \cdot p(x) - f^*(T(x)) \cdot q(x) \, dx \\ &\geq \sup_{\hat{T} \in \mathcal{T}} \int_{\mathcal{X}} \hat{T}(x) \cdot p(x) - f^*(\hat{T}(x)) \cdot q(x) \, dx \\ &= \sup_{\hat{T} \in \mathcal{T}} E_{x \sim q} [\hat{T}(x)] - E_{x \sim q} [f^*(\hat{T}(x))] \end{aligned}$$

lower bound is likelihood-free w.r.t p and q

$$f^*(y) = \sup_{x \in \text{dom} f} (yx - f(x))$$

$$f^{**}(x) = \sup_{y \in \text{dom} f^*} (xy - f^*(y))$$

$$T(x) = \arg \max_{y \in \text{dom} f^*} y \cdot \frac{p(x)}{q(x)} - f^*(y)$$

restricting $\hat{T}$ to a family

let $P = P_{\text{data}}$, $Q = P_0$, Parametrize $\mathcal{T}$ by ϕ and $\mathcal{G}$ by θ

$$\min_{\theta} \max_{\phi} \left\{ F(\theta, \phi) = E_{x \sim P_{\text{data}}} [T_{\phi}(x)] - E_{x \sim P_0} [f^*(T_{\phi}(x))] \right\}$$

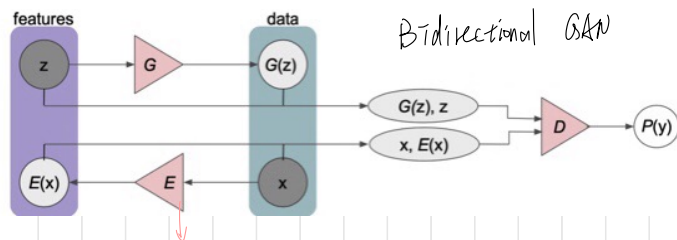
max over T_{ϕ} to approximate $D_f(P_{\text{data}}, P_0) \approx E_{x \sim P_{\text{data}}} [T(x)] - E_{x \sim q} [f^*(T(x))]$
min over G_{θ} to minimize $D_f(P_{\text{data}}, P_0)$

} Choose a D_f , compute the f^*
all fit into the GAN scheme

· Inferring Latent Representations in GAN : BiGAN

1. For any sample x , use the activations of pre-final layer of discriminator as a feature representation
2. directly infer the latent variable z of the generator

for any (x, z) pair sampled from generator, can evaluate $P(z)$



the encoder only observe $x \sim P_{\text{real}}(x)$ during training
to learn a mapping $E: x \rightarrow z$

at test time, new samples are generated via G
latent representation are inferred via E

CycleGAN: Adversarial Training across Two domains

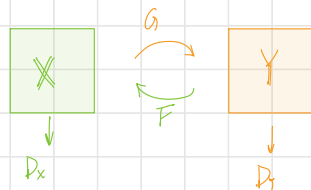
To match two distributions, learn two parameterized conditional generative model $G: X \rightarrow Y$ $F: Y \rightarrow X$

$G: X \rightarrow Y$ maps an element of X to an element of Y

A discriminator D_Y compares observed dataset Y and generated samples $\hat{Y} = G(X)$

$F: Y \rightarrow X$ maps an element of Y to an element of X

A discriminator D_X compares observed dataset X and generated samples $\hat{X} = F(Y)$



problem: G ignores X and just try to generate a $\hat{Y}$
such that $D_Y(\hat{Y}) = 1$

cycle consistency: if can go from X to $\hat{Y}$ from G ,
then it should be possible to go from $\hat{Y}$ back to X via F

$$\left. \begin{array}{l} F(G(X)) \approx X \\ G(F(Y)) \approx Y \end{array} \right\}$$

$$\min_{F, G, D_X, D_Y} \mathcal{L}_{GAN}(G, D_Y, X, Y) + \mathcal{L}_{GAN}(F, D_X, X, Y) + \lambda \left[E_X \|F(G(X)) - X\| + E_Y \|G(F(Y)) - Y\| \right] \quad \left. \vphantom{\min_{F, G, D_X, D_Y}} \right\} \text{cycle consistency}$$

when G is a normalizing Flow Model

$F = G^{-1}$ exact cycle-consistent

need only one transformation

can train via MLE and/or adversarial training

