

GANs: Generative Adversarial Networks

Goodfellow et al (2014)

Idea: adversarial training min-max objective

was SOTA

on images
until recently
(now diffusion)

still some advantages
— computationally
lighter & faster

saddle pt
not global min!

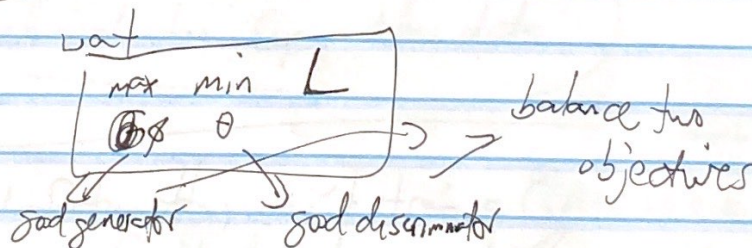
train 2 NNS
to go to defeat one another

G generator $z \rightarrow x$ (noise data)
 D discriminator
classifies generated "fake"
vs data "real"

GAN loss:

BCE for D
but also loss for G !

$$L(\theta) = - \left(\sum_{i \in \text{data}} \log D(x_i) + \sum_{i \in \text{gen}} \log (1 - D(G(z_i))) \right)$$



Why does this work??

— formally: $\min_{\theta} L \rightarrow D_{\theta}(x) = p_{\text{data}}(x)$ (Neyman-Pearson Lemma!)

← plug back into L

$$L|_{D=D_{\theta}} = - \left(\sum_{x \in \text{data}} \log \frac{p_{\text{data}}(x)}{p_{\text{data}}(x) + p_{G_{\theta}}(x)} + \sum_{x \in \text{gen}} \log \frac{p_{G_{\theta}}(x)}{p_{\text{data}}(x) + p_{G_{\theta}}(x)} \right)$$

$$\rightarrow - \int dx \left(p_{data} \log \frac{p_{data}}{p_{data} + p_G} + p_G \log \frac{p_G}{p_{data} + p_G} \right)$$

$$= -2 \text{JSD}[p_{data}, p_G] + \text{const}$$

↓
 Jensen-Shannon Divergence — distance
 $0 \leq \text{JSD}^{(pq)} \leq 1$ btw prob
 zero iff $P=Q$ dist's!

$$\max_{\beta} L \Big|_{\beta = \beta_g} = \min_{\beta} \left(\text{JSD}[p_{data}, p_{G_{\beta}}] + \text{const} \right)$$

$$\Rightarrow \boxed{p_{G_{\beta}} = p_{data}}$$

(can also
 prove more
 straightforwardly
 from integral
 above)

so formally, GAN objective is
 achieved when G_{β} matches data!

↓
 key pt:
 p_{data} only
 learned
 implicitly!

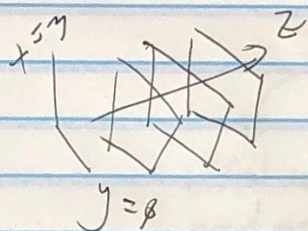
Unfortunately, can't train a GAN so easily,
 so formal derivation not 100% applicable...

→ go to GAN MNIST Demo

→ Training:
 Freeze G
 Train D (batch real + batch gen)
 Freeze D
 Train G (batch gen)
 repeat

CALOGAN example (1705.02355, 1712.10321)

- Early example of ^{deep} gen model for fast sim - surrogate modeling (first in HEP?)
- Used vanilla GAN for GEANT4 e^+, γ, π^+ calo showers by ATLAS ~~ECAL~~ ECAL



3 layers of LAr + lead absorber
simplification: turn sample fraction off
(100% efficient)

3 GANs: e^+, γ, π^+

Data: $\begin{cases} 1-100 \text{ GeV unif } E_{inc} \\ 100k \text{ showers each generated w/ GEANT4} \end{cases}$

"Voxelization"
 $3 \times 96 = 12 \times 12 = 12 \times 6$
 $= 504 \text{ total voxels}$

energy in each voxel
 $\vec{E} \in \mathbb{R}^{504}$

↓
to generate from
goal: learn $p(\vec{E} | E_{inc})$

→ to overview of paper
main results: quality so-so

speed - very fast (but CPU/GPU
comparison subtleties)
metrics?

Problems w/ (Vanilla) GANs

- Convergence: min-max unstable

vanishing gradients

if D too powerful, can overwhelm G
(100% separable, no useful gradients for G)
can't train D to completion

if D too weak, random guessing
also no gradients for G

GAN never converges
good performance quickly destabilized

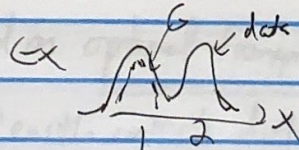
- Model selection

- D & G losses not correlated w/ quality
- in industry, / natural images
often rely on "by eye" test!
- hard to know when to stop training

- Mode collapse

G & D stuck in vicious cycle

gan lab demo



G learns to produce 1's

D fooled for a while

then realizes 2 are all real, guesses on 2's

→ G will switch to making 2's

Beyond vanilla GANs

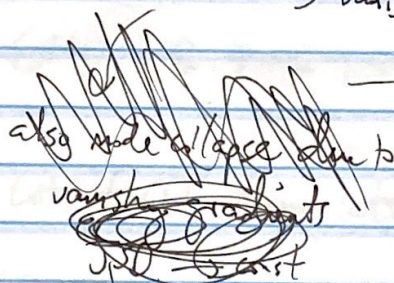
- Many issues w/ GANs stem from loss fn

D trained to optimality $\rightarrow \min_G JSD(p_{data}, p_G)$

$\underbrace{p_{data} \leftrightarrow p_G}_{JSD=1}$

$JSD=1$ whenever p_{data} & p_G disjoint
bounded from above
saturates

$\rightarrow$ Vanishing gradients!



$\rightarrow$ when G is very bad $\rightarrow$ no ability to improve!

So can't train D to optimality for vanilla GAN

also causes mode collapse, stems from suboptimal D .

- Want: better GAN Loss, unbounded from above
informative even for widely separated dists

$\rightarrow$ one SOTA approach: Wasserstein GANs (1701.07875
Arjovsky, Chintala, Bottou)

Based on optimal transport theory

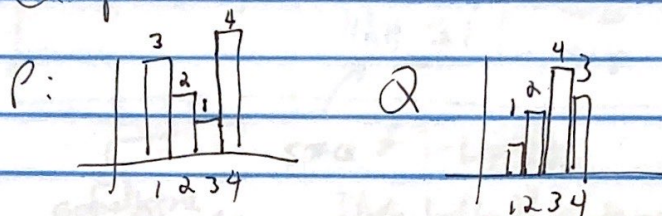
dist P to
dist Q
 P, Q

"earth mover's distance" aka "Wasserstein distance"

"work it takes to transform P into Q "

$\hookrightarrow$ right property: more separated P & Q ,
more work

Example:



- transport 2 units $1 \rightarrow 2$ so 1 matches
- transport 2 units $2 \rightarrow 3$ so 2 matches
- transport 1 unit $4 \rightarrow 3$ so 3 & 4 match ✓

$$\text{EMD: } 2 \cdot 1 + 2 \cdot 1 + 1 \cdot 1 = 5.$$

↳ Lots of OT theory later...

like MSE
in prob
dist'n
space

$$\text{EMD}(P, Q) = W(P, Q) = \min_{\gamma \in \Pi(P, Q)} \mathbb{E}_{(x, y) \sim \gamma} [\|x - y\|]$$

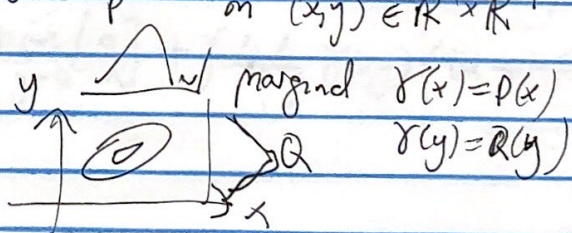
arb. dist'n
 $P(x), Q(y), x, y \in \mathbb{R}^d$

"Wasserstein-1"
✓
Wasserstein-k
is $\sqrt[k]{\mathbb{E} \|x - y\|^k}$

expectation value
over all dist'n's γ
coupling
distance

space of all prob dist'n's $\Pi(x, y)$
on $(x, y) \in \mathbb{R}^d \times \mathbb{R}^d$

$W_1 = 0 \Leftrightarrow P = Q$
unbounded for $P \neq Q$
 $\uparrow \Leftrightarrow \downarrow$ JSO $\rightarrow 1$
 $W_1 \rightarrow \infty$



W_1 totally intractable...

→ Brilliant idea: use Kantorovich-Rubinstein duality

Amazing formula

K-R duality

$$W_1(P, Q) = \max_{\|h\|_L \leq 1} \left[\mathbb{E}_{x \sim p} [h(x)] - \mathbb{E}_{y \sim q} [h(y)] \right]$$

↙
actually true
for K-Lipschitz
 $\|h\|_L \leq K$

space of 1-Lipschitz

$$|h(x_1) - h(x_2)| \leq |x_1 - x_2| \quad \forall x_1, x_2$$

$$\Downarrow$$

$$|h'(x)| \leq 1 \text{ infinitesimal } \int \text{ by mean value theorem}$$

↘ rough idea of pf: ^{comes from} ~~rough~~ Lagrange multipliers

$$\left\{ \begin{aligned} \mathbb{E}_{x \sim p} (|x - y|) &= 0 \quad \text{at } x=y \\ &+ \int dx dy \delta(x, y) (|x - y|) \\ &+ \int (p(x) - \int \delta(x, y) dy) f(x) dx \\ &+ \int (q(y) - \int \delta(x, y) dx) g(y) dy \\ &= \mathbb{E}_{x \sim p} [f] + \mathbb{E}_{y \sim q} [g] + \int dx dy \delta(x, y) (|x - y| - f(x) - g(y)) \\ &\vdots \end{aligned} \right.$$

w/ K-R duality, W_1 becomes tractable!

Approx h w/ NN $\rightarrow$ "critic"
analog of discriminator

$\max_{\|h\|_L \leq 1} (\dots) \rightarrow$ training the critic!

WGAN loss

$$L = \sum_{x \in \text{real}} h(x) - \sum_{z \in \text{latent}} h(G(z))$$

$$\min_G \max_{h \in \text{Lipschitz}} L$$

Enforcing Lipschitz: many approaches

• original paper: "weight clipping"

after each weight update

$$w \rightarrow \text{clip}(w, -c, c) \quad \begin{cases} \text{if } w > c \\ w = c \\ \text{if } w < -c \\ w = -c \end{cases}$$

This enforces K -Lipschitz crudely

roughly: NN output $h = g_1 \circ w_1 \circ g_2 \circ w_2 \circ \dots \circ g_L \circ w_L(x)$

$$\text{So } \|\nabla_x h\| \sim \|\partial_{w_1} h, g_2' w_2, \dots, g_L' w_L\|$$

So if weights bounded, $\|\nabla_x h\|$ bounded
& activation gradients g_i' then bounded

• Better way: gradient penalty

$$\text{add regularizer to loss } \mathcal{L} = \mathbb{E}_{\hat{x} \sim p} \left(\|\nabla_x h(\hat{x})\| - 1 \right)^2$$

might not
be state of
the art...

p : dist'n of pts $\hat{x} = tx + (1-t)y$
 $x, y \sim \text{data, gen}$
 $t \sim U(0, 1)$

Some theory behind this...

not clear why seems to require $\|\nabla_x h\| = 1$
instead of ≤ 1 ... but empirically works