

Generative Models

Claudius Krause

@ fh.phys.uni-heidelberg
.de

based on "Modern Machine Learning in LHC Physics"
by T. Plehn, A. Butter, B. Dillou, C. Krause
arxiv: soon

Rutgers Physics 694 physics.rutgers.edu/~dshih/694

Generative Models

problem: have a distribution $p_{\text{truth}}(x)$

I want to sample ("generate") new

$$x \sim p_{\text{truth}}(x)$$

p can be given explicitly as function $\frac{d\phi}{d\theta}$

- implicitly via a set of data points

$\{x\}$ defining $p_{\text{data}}(x) (\approx p_{\text{truth}}(x))$

applications in HEP

- event generation (partonic, end-to-end)
- phase space
- calorimetric showers
-

this is a stochastic (random) process (RNG) $\rightarrow$ need random input

$$z \sim p_{\text{ident}}(z) \quad \Rightarrow \quad x = f_{\theta}(z) \sim p_{\text{model}}(x) \approx p_{\text{data}}(x)$$

$\hookrightarrow$ easy: Gaussian, uniform $\approx p_{\text{truth}}(x)$

these lectures: how to construct & train $f_{\theta}(z)$

VAEs

GANs

Normalizing Flows

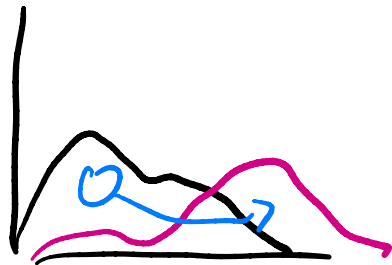
How do we know, we have a good $f_\theta(z)$?

if we have $p_{\text{model}} \approx p_{\text{data/true}}$ explicitly : Kullback-Leibler (KL) divergence

$$\begin{aligned} 1) D_{KL} [p_{\text{data}}, p_{\text{model}}] &= \int dx \, p_{\text{data}}(x) \log \frac{p_{\text{data}}}{p_{\text{model}}} \\ &= \langle \log \frac{p_{\text{data}}}{p_{\text{model}}} \rangle_{\text{data}} \end{aligned}$$

2) Earth Mover's Distance (Wasserstein metric)

works for continuous (p known)
or discrete (only samples available)



discrete p

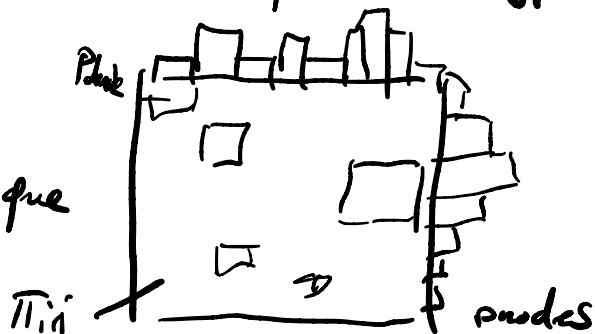
$$p_{data,i} \\ i \in [1, N_1]$$

$$p_{model,j} \\ j \in [1, N_2]$$

define transport strategy π_{ij}

as a matrix relating the
two sets

π_{ij} not unique



$$p_{data,i} = \frac{1}{N_2} \sum_j \pi_{ij}$$

$$p_{model,j} = \frac{1}{N_1} \sum_i \pi_{ij}$$

$$W[P_{data}, P_{model}] = \min_{\pi} \frac{1}{N_1 N_2} \sum \pi_{ij} |x_i - y_j|$$

$\Rightarrow$ scalar horribly w/ dimensionality "amount of dirt" $\times$ distance

$\mathcal{L}$ is expensive to compute

3) generate events & look at single events (does a face look realistic?)
 look at distributions $\hookrightarrow$ (dataset dep.)
 $\hookrightarrow$ (always a projection)

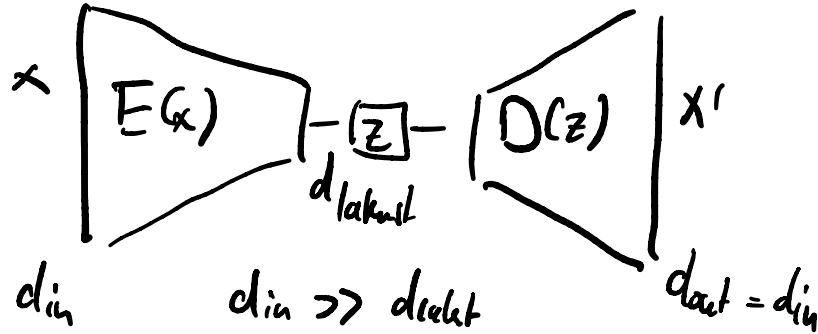
4) Classifier test [1610.06545] $\Rightarrow$ train NN to distinguish data vs. gen. samples

if $P_{\text{correct}} = 0.5$ (classifier is confused)

$$\hookrightarrow P_{\text{data}}(x) = P_{\text{model}}(x)$$

1 / Variational Auto-encoder VAE 13.12.6.11.4

start w/ AE part first



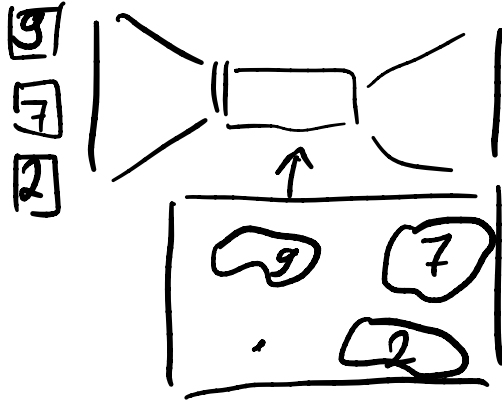
loss: MSE

$$\mathcal{L}_{MSE} = \frac{1}{N} \sum_{i=1}^N \|x_i - x'_i\|^2$$
$$x'_i = D(E(x_i))$$

- compress information
("lossy") $\rightarrow$ will learn most useful latent representation of data

- generative model: sampling from z & use $D(z)$
But: latent space not regularized

- what distribution to sample from?
- interpolation $\rightarrow$ gaps



$$z \sim E_{\theta}(z|x)$$

$$x' \sim D_{\theta}(x'|z)$$

VAE give structure to latent space

$$x \left| \begin{array}{c} \text{E} \end{array} \right| \left\{ \begin{array}{c} \mu \\ \sigma \end{array} \right\} z = \mu + \sigma \cdot \varepsilon$$

$\varepsilon \sim N(0, 1)$

$$\left| \begin{array}{c} \text{D} \end{array} \right| x'$$

given data x , we want to know / infer z of the gen. model

$$D(z|x) = \frac{D(x|z) \cdot p_z(z)}{p_{\text{data}}(x)}$$

VAE : approximate $D(z|x)$ by $E(z|x)$ (variational inference)

$$D_{KL}[E(z|x), D(z|x)] = \left\langle \log \frac{E(z|x)}{D(z|x)} \right\rangle_{E(z|x)} =$$

$$= \left\langle \log E(z|x) - \log \frac{D(x|z) p_z(z)}{p_{\text{data}}(x)} \right\rangle_{E(z|x)}$$

$$= \underbrace{\left\langle -\log D(x|z) \right\rangle_{E(z|x)} + D_{KL}[E(z|x), p_z(z)]}_{\text{ELBO} \rightarrow \text{bounds LL}} + \log p_{\text{data}}$$

$$\equiv \mathcal{L}_{MSE} + \beta_{KL} D_{KL} [E(z|x), N(0,1)]$$

↓
if Gaussian dist: MSE reconstruction

regulates the latent space to be Gaussian

latent space



1 1 1 1 1 8

Sampling: from $N(0,1)$ & then passing through $D(z)$

for practical applications : we need a conditional VAE

- condition (discrete or continuous) needs to be given to $E \& D$

$D(z, c)$ generates x for a given c

pros:

- easy to setup / train
- explicit map to the latent space

cons

- rarely state of the art
- no optimality guarantee

[Z] Generative Adversarial Networks (GAN) 1406.2661

train 2 adversarial NN:

$$\text{generator } g : \underset{\substack{\downarrow \\ \text{init. / gauss}}}{z} \rightarrow x$$

$$\underset{\substack{\text{discriminator} \\ \text{critic}}}{D} \left(\underset{x}{\cdot} \underset{g(z)}{\cdot} \right) = \begin{cases} 0 & \text{generated data} \\ 1 & \text{true data} \end{cases}$$

train D by minimizing $\langle 1 - D(x) \rangle_{x \sim p_{\text{true}}} \quad \& \quad \langle D(x) \rangle_{x \sim p_{\text{gen}}}$

enhance sensitivity $1 - 0 \rightarrow -\log D$

$D \rightarrow -\log(1 - D)$

$$\mathcal{L}_D = \langle -\log D \rangle_{p_{\text{true}}} + \langle -\log(1 - D) \rangle_{p_{\text{gen}}} \quad \text{BCE}$$

train g by making it generate samples that look like real data

$$L_g = \langle -\log D \rangle_{p_{\text{gen}}} \sim \langle \log(1-D) \rangle_{p_{\text{gen}}}$$

Same
minimum, better gradients

together:

$$\min_g \max_D \langle \log D(x) \rangle_{p_{\text{real}}} + \langle \log(1-D(g(z))) \rangle_{p_z}$$

$\rightarrow$ is a good classifier
 $\hookrightarrow$ looks D

Optimal point: Nash equilibrium

$$D_{opt} = \frac{P_{true/data}}{P_{true/data} + P_{gen}}$$

$$\mathcal{L}_D \rightarrow - \left\langle \log \frac{P_{data}}{P_{data} + P_{gen}} \right\rangle_{P_{data}} - \left\langle \log \frac{P_{gen}}{P_{data} + P_{gen}} \right\rangle_{P_{gen}}$$

$$= - D_{KL}(P_{data}, P_{data} + P_{gen}) - D_{KL}(P_{gen}, P_{data} + P_{gen})$$

$$= - 2 D_{JS}[P_{data}, P_{gen}] + 2 \log 2$$

$$\mathcal{L}_g \rightarrow \left\langle \log (1 - D_{opt}) \right\rangle_{P_{gen}}$$

$$= D_{KL}\left[P_{gen}, \frac{P_{data} + P_{gen}}{2}\right] - \log 2$$

$\hookrightarrow$ minimal for $P_{gen} = P_{data}$

training :

here G

take a batch from training data
- " - gen data } train D

repeat

here D

take sample from generated data $\rightarrow G$

problems w/ GANs

- convergence min-max unstable

vanishing gradients from D

↳ by regularization

1705.09367

1801.04406

- metrics for success

- mode collapse

GAN gets stuck in vicious cycle



G tries to produce realistic 1

D will be fooled

then D will learn 2's are real

then G will focus on 2's

ways to improve GANs

→ replace discriminator by critic

① learn Wasserstein distance between
 P_{data} & P_{gen}

WGAN 1701.07875

GAN

pro: great, realistic single images

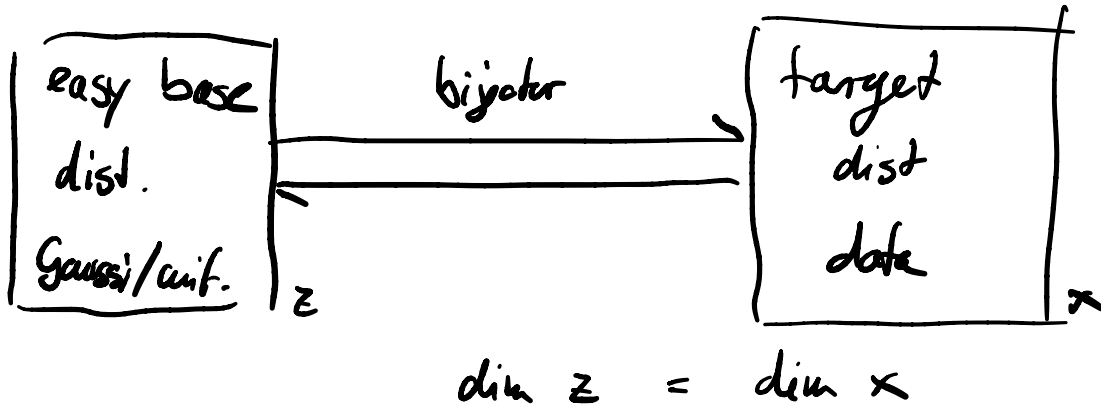
con: distributions only @ $\Theta(10\%)$

3 | Normalizing Flows

VAEs & GANs learn p implicitly (only access to samples, not $p(x)$)

Flows learn $p(x)$ explicitly

key idea: learn complicated coordinate transformation



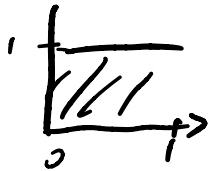
for coordinate transform $\vec{x} = f(\vec{z})$

$$p(x) = \pi(z) \left| \det \frac{df(z)}{dz} \right|^{-1}$$
$$= \pi(f^{-1}(x)) \left| \det \frac{\partial f^{-1}(x)}{\partial x} \right|$$

→ need jacobian determinant

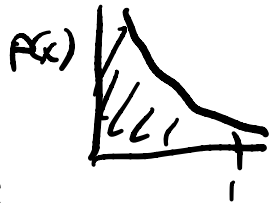
→ need inverse $\vec{z} = f^{-1}(\vec{x})$

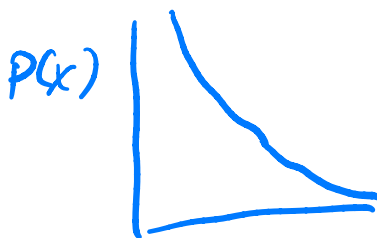
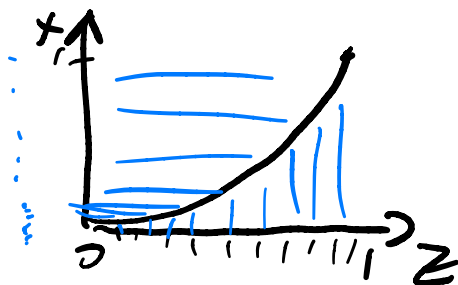
$\pi(z)$ uniform



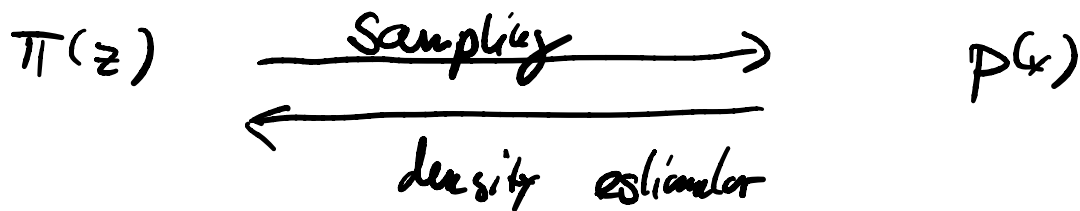
$$x = f(z) = z^2$$

$$p(x) = 1 \cdot |2z|^{-1} = \frac{1}{2\sqrt{x}}$$





bijector works in both ways

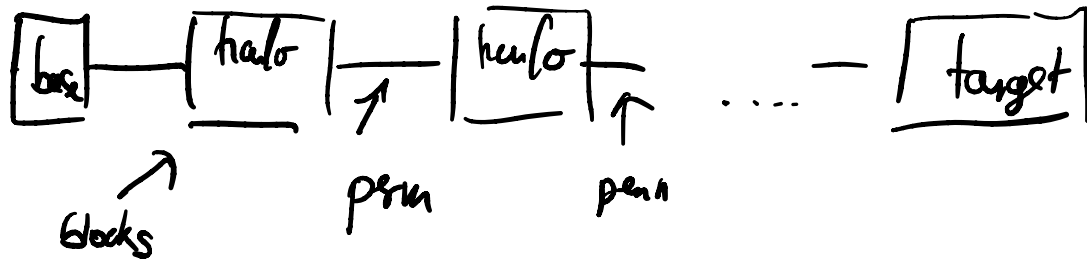


first idea: $f(z) = \text{NN}(z)$

does not work well

- invert a NN?
- Jacobian $\sim \mathcal{O}(n^3)$

Flows avoid these problems by learning
 parameters of a series of invertible, pre-defined
 transformations (w/ hackable jacobian)



chain of
 bijectors
 is a bijector

1505.05770

$$x = C_K \circ C_{K-1} \circ C_{K-2} \dots C_1(z_0) \quad z_0 \sim \pi_0(z_0)$$

$$p(x) = \pi_0(z_0) \prod_{k=1}^K \left| \det \frac{\partial C_k}{\partial z_{k-1}} \right|^{-1}$$

training - max $\log p$ (or min. NLL)

- use D_{KL} to fit to (normalized) target function

$\Rightarrow$ more stable than GANs, better performance than VAEs

in more detail

inside each block, assume the halo functions

$$\vec{c}(\vec{z}, \vec{\mu}) = \left(c_1(z_1, \vec{\mu}_1), c_2(z_2, \vec{\mu}_2), \dots, c_n(z_n, \vec{\mu}_n) \right)^T$$

$\vec{\mu}_i$ still depend on z_j ($i \neq j$) to
capture correlations

there are 2 main architectures to tame Jacobians

↳ $\mathcal{O}(n^3) \Rightarrow \mathcal{O}(n)$
matrix is triangular

① autoregression models

$$\vec{\mu}_1 = \text{const.}$$

$$\downarrow$$

$$p(x_1)$$

$$\vec{\mu}_2(z_1)$$

$$\nwarrow$$

$$p(x_2 | x_1)$$

$$\vec{\mu}_3 = \vec{\mu}_3(z_1, z_2)$$

$$\swarrow$$

$$p(x_3 | x_2, x_1) \quad \dots$$

$$\vec{\mu}_i(z_1, \dots, z_{i-1})$$

$$p(x) = \prod_{i=1}^n p(x_i | x_1, \dots, x_{i-1})$$

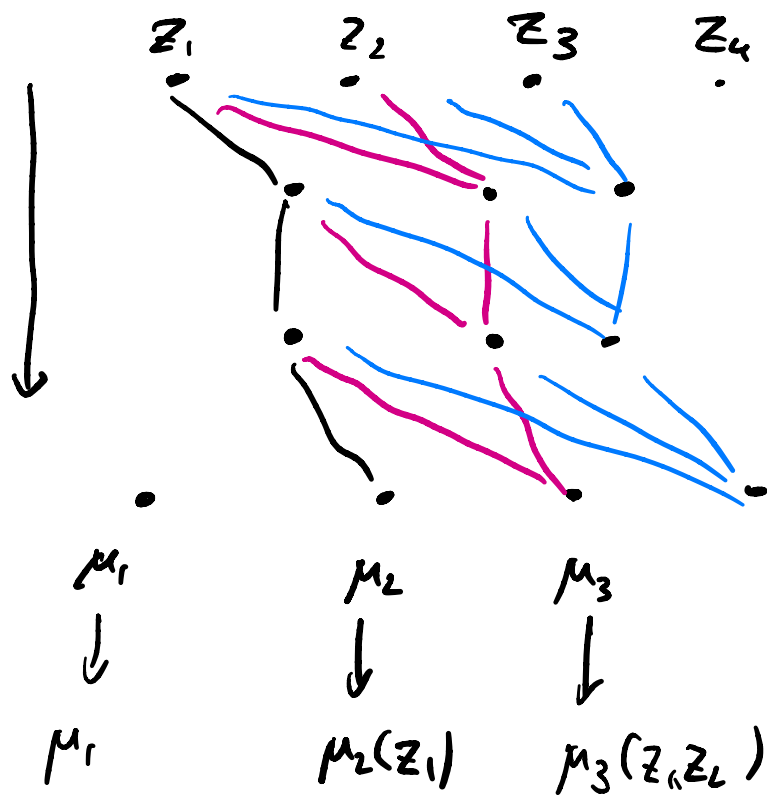
$$x_i = c_i(z_{1:i})$$

$$J = \begin{pmatrix} \triangle & 0 \end{pmatrix} \quad \mathcal{O}(n)$$

$$\det J = \prod_{i=1}^n J_{ii}$$

NN realization : MADE (Masked Autoregressive for Density Estimation)

1502.03503



forward: single call gives all μ_s fast

inverse (loop n-times slow to get all z_s)

fast	inferre $p(x)$	Sampling
slow	Sampling	inferre $p(x)$

MAF

1705.07057

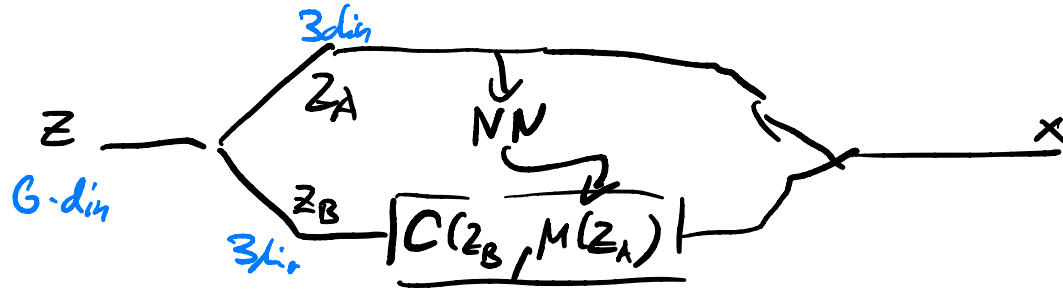
IAF

1606.04934

② Coupling Layer based flows (bipartite) INN

Real NVP 1605-08803

Split z in 2 sets z_A, z_B



forward

$$x_A = z_A$$

equally fast in both
directions

$$x_B = C_B(z_B, \mu(z_A))$$

inverse

$$z_A = x_A$$

$$z_B = C^{-1}(x_B, \mu(z_A)) = x_B$$

$$(\det J) = \begin{vmatrix} 1 & 0 \\ \frac{\partial C_B}{\partial z_A} & \frac{\partial C_B}{\partial z_B} \end{vmatrix}$$

diagonal

how about $C(z, \vec{\mu})$

↳ detailed choice depends on domain

base gaussian	$(-\infty, \infty)$	} logit sigmoid
uniform	$(0, 1)$	

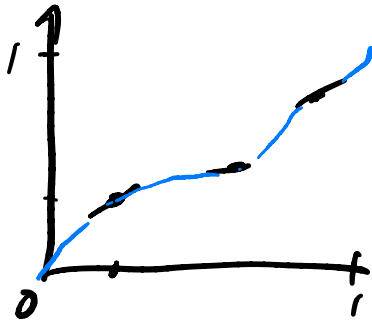
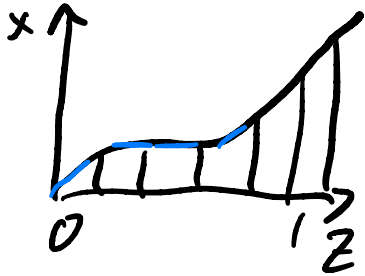
affine coupling layer

$$C(z, \underset{\substack{\vec{s}, \vec{t}}}{\vec{\mu}}) = z \odot \exp(\vec{s}) + \vec{t} = x$$

$$\text{inuse} : (x-t) \circ \exp(-s) = z$$

$$|\det J| = \exp(\sum s)$$

piecewise splines (monotonic)



- linear, quadratic 1808.03856

- Rational Quadratic Splines 1906.04032

$$C = \frac{a_2 \alpha^2 + a_1 \alpha + a_0}{b_2 \alpha^2 + b_1 \alpha + b_0}$$

α ... location inside bin
 a_i, b_i given by NN