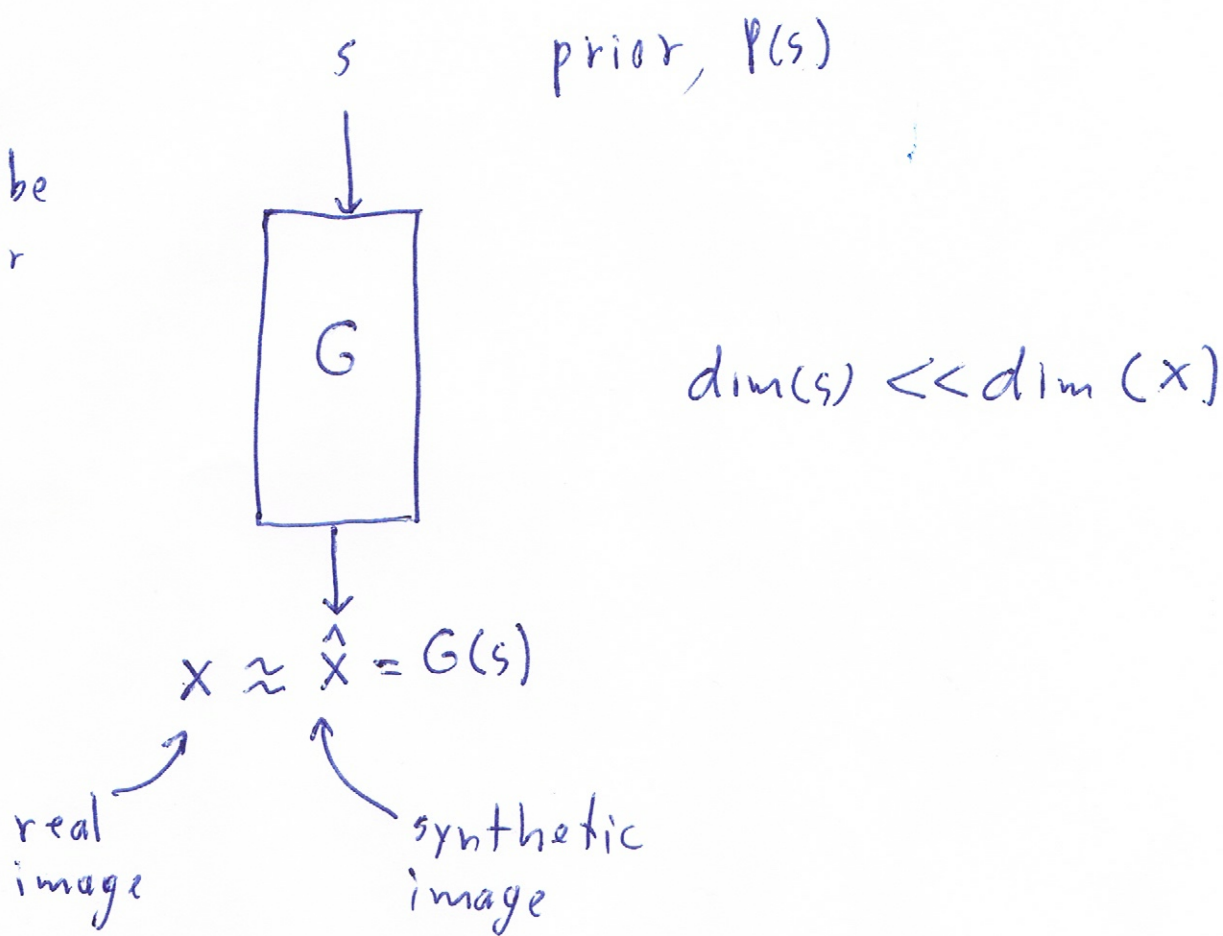


Generative Models

①

G may be given or learned.



Pretend that we can randomly generate images "near" $\hat{x}$ with probability $P(x|\hat{x})$

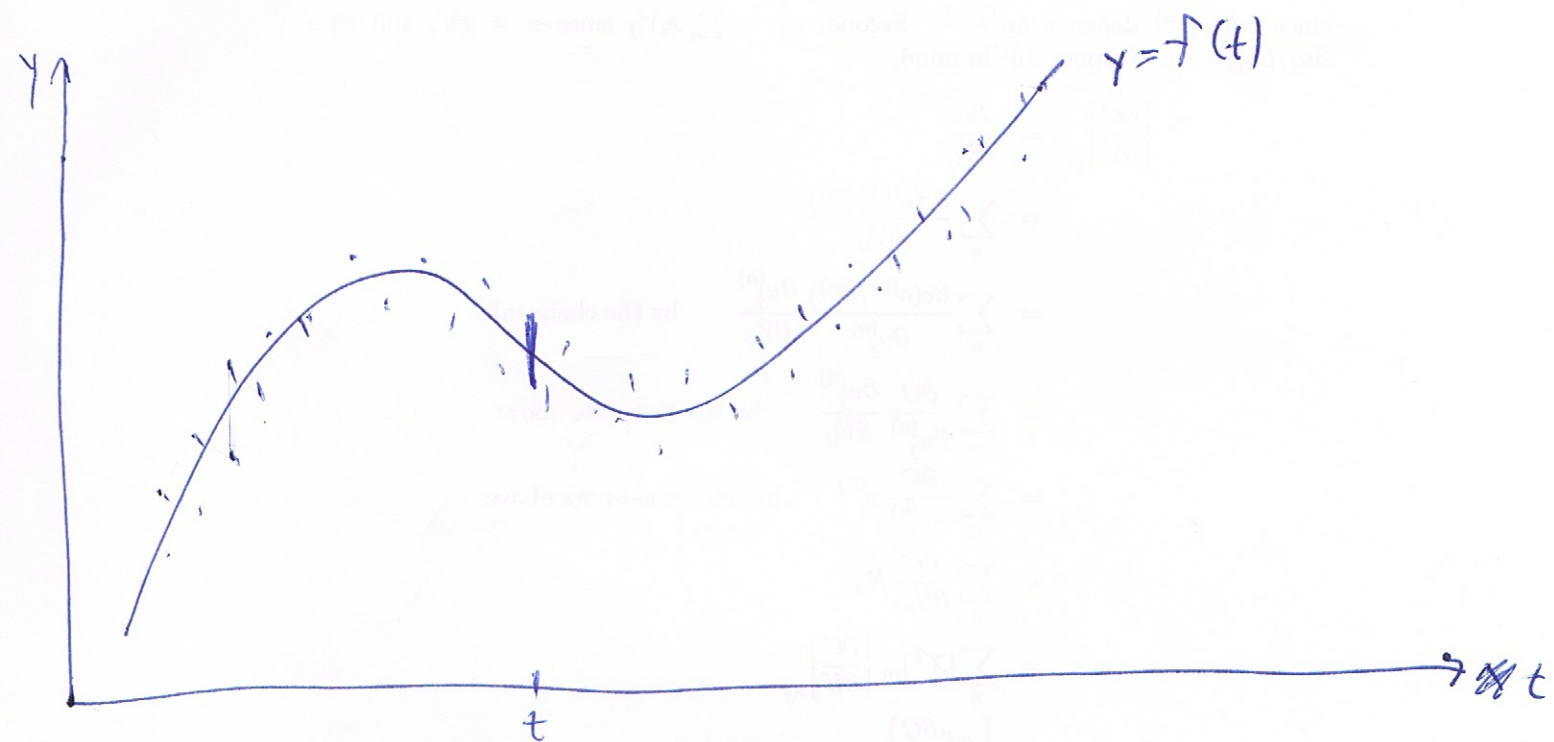
$$\text{eg. } P(x|\hat{x}) = \mathcal{N}(x|\hat{x}, \sigma) = \frac{e^{-\|x - \hat{x}\|^2 / 2\sigma^2}}{(\sqrt{2\pi}\sigma)^k}$$

$$k = \dim(x)$$

so, images similar to $\hat{x}$ are most likely.

(2)

compare to regression



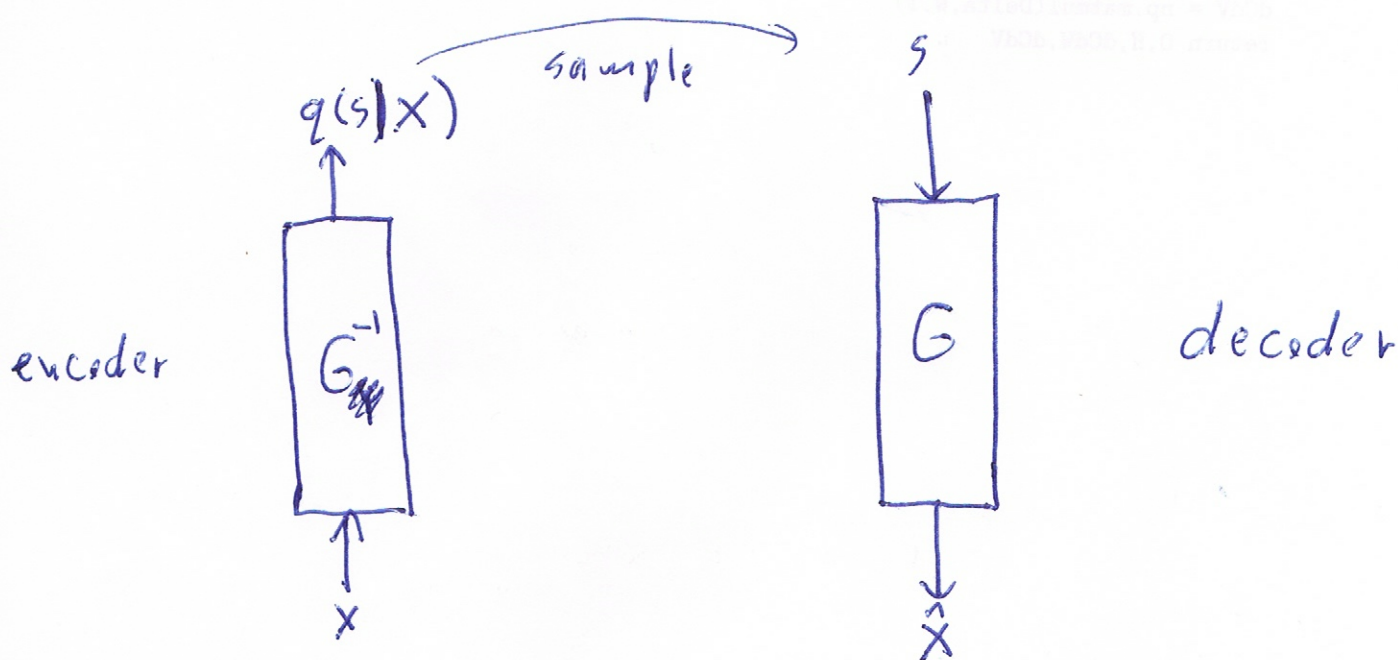
$$P(y|t) = \mathcal{N}(y|f(t), \sigma)$$

ie, given t , we predict a distribution for y . Typically, $f(t)$ is the mean of the distribution. In simple models, σ is fixed, independent of t .

Inverse Graphics

(3)

- compute (estimate) $P(s|x)$
- intractable in general.
- approximate $P(s|x)$ using a NN (encoder).



Train G^{-1} (& maybe G)

so that $x \approx \hat{x}$ for all training data.

Backprop through G & G^{-1} . How?

Variational Autoencoders

Given training data $x^1 \dots x^n$,
train so that $P(x^1, \dots, x^n)$ is maximal.

$$P(x) = \sum_s P(x, s) = \sum_s P(x|s) \cdot P(s)$$

$\uparrow$
 huge
sum

$\uparrow$
 known
learnable

$\uparrow$
 known

Assume iid data, so $P(x^1 \dots x^n) = \prod_n P(x^n)$

$$\therefore \log P(x^1 \dots x^n)$$

$$= \sum_n \log P(x^n)$$

$$= \sum_n \log \left\{ \sum_s P(x^n, s) \right\} \quad \left. \begin{array}{l} \text{sum inside} \\ \log \end{array} \right\}$$

$$= \sum_n \log \left[\sum_s q(s|x^n) \cdot \frac{P(x^n, s)}{q(s|x^n)} \right]$$

(5)

$$> \sum_n \sum_s q(s|x^n) \log \frac{p(x^n, s)}{q(s|x^n)} \stackrel{\text{def}}{=} \mathcal{L}$$

- true for any distribution, q ,
because $\log$ is concave.

Q.E.D.

exercise:

$$\mathcal{L} = \log p(x^1 \dots x^n)$$

$$- \sum_n \text{KL}[q(s|x^n), p(s|x^n)]$$

$\therefore$ maximizing $\mathcal{L}$ minimizes ^{KL, i.e.,} the
difference between $q(s|x^n)$ & $p(s|x^n)$
i.e., makes $q(s|x^n)$ a better
approximation of $p(s|x^n)$.

Problem

(S.1)

The sum over S is huge
& can't be evaluated. But we
can approximate (to arbitrary accuracy):

$$\mathcal{L} = \sum_{n=1}^N E_{s \sim q} \log \frac{P(x^n, s)}{q(s|x^n)} \quad *$$

$$\approx \sum_{n=1}^N \frac{1}{L} \sum_{\ell=1}^L \log \frac{P(x^n, s_\ell^n)}{q(s_\ell^n|x^n)}$$

where $s_\ell^n \sim q(s|x^n)$

If N is large, can use $L=1$.

This works because $q(s|x^n)$ is
relatively narrow.

(6)

$$\therefore \mathcal{L} \approx \sum_n \log \frac{P(X^n | s^n)}{q(s^n | X^n)}$$

where $s^n \sim q(s^n | X^n)$

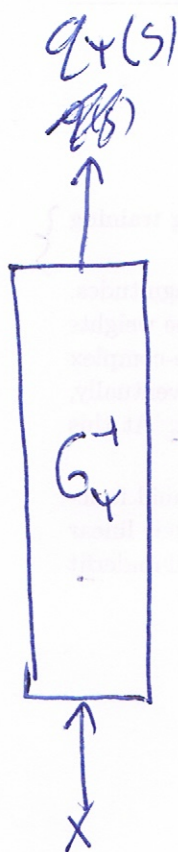
$$= \sum_n \log \frac{P(X^n | s^n) \cdot P(s^n)}{q(s^n | X^n)}$$

$$= \sum_n [\log P_\theta(X^n | s^n) + \log P(s^n) - \log q_\psi(s^n | X^n)]$$

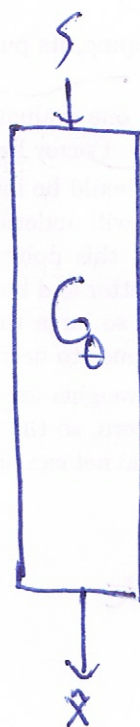
Note: $\nabla_\theta \mathcal{L} = \sum_n \nabla_\theta \log P_\theta(X^n | s^n)$

⑥ ⑦

parameters
 ψ



parameters
 θ



$p_\theta(x|s)$

$$\begin{aligned}
 p_\theta(x|s) &= p(x | \hat{x} = G_\theta(s)) \\
 &= \mathcal{N}(x | G_\theta(s), \alpha) \\
 &= \frac{e^{-(x - G_\theta(s))^2 / 2\alpha}}{(\sqrt{2\pi}\alpha)^k}
 \end{aligned}$$

$k = \dim(x)$

$$\begin{aligned}
 \log p_\theta(x|s) &= -\frac{(x - G_\theta(s))^2}{2\alpha} \\
 &\quad - k \log(\sqrt{2\pi}\alpha)
 \end{aligned}$$

$$\nabla_{\theta} \mathcal{L} = \sum_n \nabla_{\theta} \log p_{\theta}(x^n | \mathcal{S}^n)$$

(9)

$$= - \sum_n \nabla_{\theta} (x - G_{\theta}(s))^2 / 2\sigma^2$$

$$= \sum_n (x - G(s)) \cdot \nabla_{\theta} G(s)$$

so, back-propagate through NN G

easy

note: G must be differentiable.

$\nabla_{\psi} \mathcal{L}$ is harder

(9)

$\mathcal{L}$ depends on the samples, ~~s, s^n~~ s^n ,
which ~~in turn depend on~~ we
draw from $q_{\psi}(s|x)$, which
depends on ψ .

Go back to *.

$$\mathcal{L} = \sum_n E_{s \sim q_{\psi}} \log \frac{p_{\theta}(x^n, s)}{q_{\psi}(s|x^n)}$$

want $\nabla_{\psi} \mathcal{L}$

$$\mathcal{L} = \sum_n E_{s \sim q_{\psi}} \log p_{\theta}(x^n, s) \} \text{ hard}$$

$$- E_{s \sim q_{\psi}} \log q_{\psi}(s|x^n) \} \text{ easy}$$

In general

(10)

$$\nabla_{\psi} E_{sq} f(s)$$

where q ~~is~~ depends on ψ .

Two approaches, depending on whether s is continuous or discrete.

① ~~continuous~~ Continuous s , (reparameterization trick).

~~example:~~ Suppose $q = \mathcal{N}(\mu, \Sigma)$

~~where μ & Σ are the outputs of a NN (typically assume Σ is diagonal)~~

~~example:~~ Suppose q is 1D
 $q = \mathcal{N}(\mu, \sigma)$

~~(easily extended to any dimension)~~

Continuous case

(11)

~~(12)~~

(reparameterization trick)

Example, Suppose s is 1D

$$q = \mathcal{N}(\mu, \sigma)$$

(easily extended to any dimension)

where μ & σ depend on ψ .

(eg, μ & σ are the output of a NN with weights ψ).

Then,

$$\mathbb{E}_{s \sim q} f(s) = \mathbb{E}_{\xi \sim \mathcal{N}(0,1)} f(\mu + \sigma \xi)$$



ie, $s = \mu + \sigma \xi$ where $\xi \sim \mathcal{N}(0,1)$

$$\therefore \nabla_{\Psi} E_{\text{avg}} f(s)$$

$$= E_{\xi \sim \mathcal{N}(0,1)} \nabla_{\Psi} f(\mu + \sigma \xi)$$

$$\nearrow E_{\xi \sim \mathcal{N}(0,1)} \left. \frac{\partial f(s)}{\partial s} \right|_{s=\mu+\sigma\xi}$$

$$= E_{\xi \sim \mathcal{N}(0,1)} \cdot f'(\mu + \sigma \xi) \nabla_{\Psi} (\mu + \sigma \xi)$$

$$= E_{\xi \sim \mathcal{N}(0,1)} f'(\mu + \sigma \xi) \left(\underbrace{\nabla_{\Psi} \mu}_{\text{vec-}\mu} + \underbrace{\xi \cdot \nabla_{\Psi} \sigma}_{\text{vec-}\sigma} \right)$$

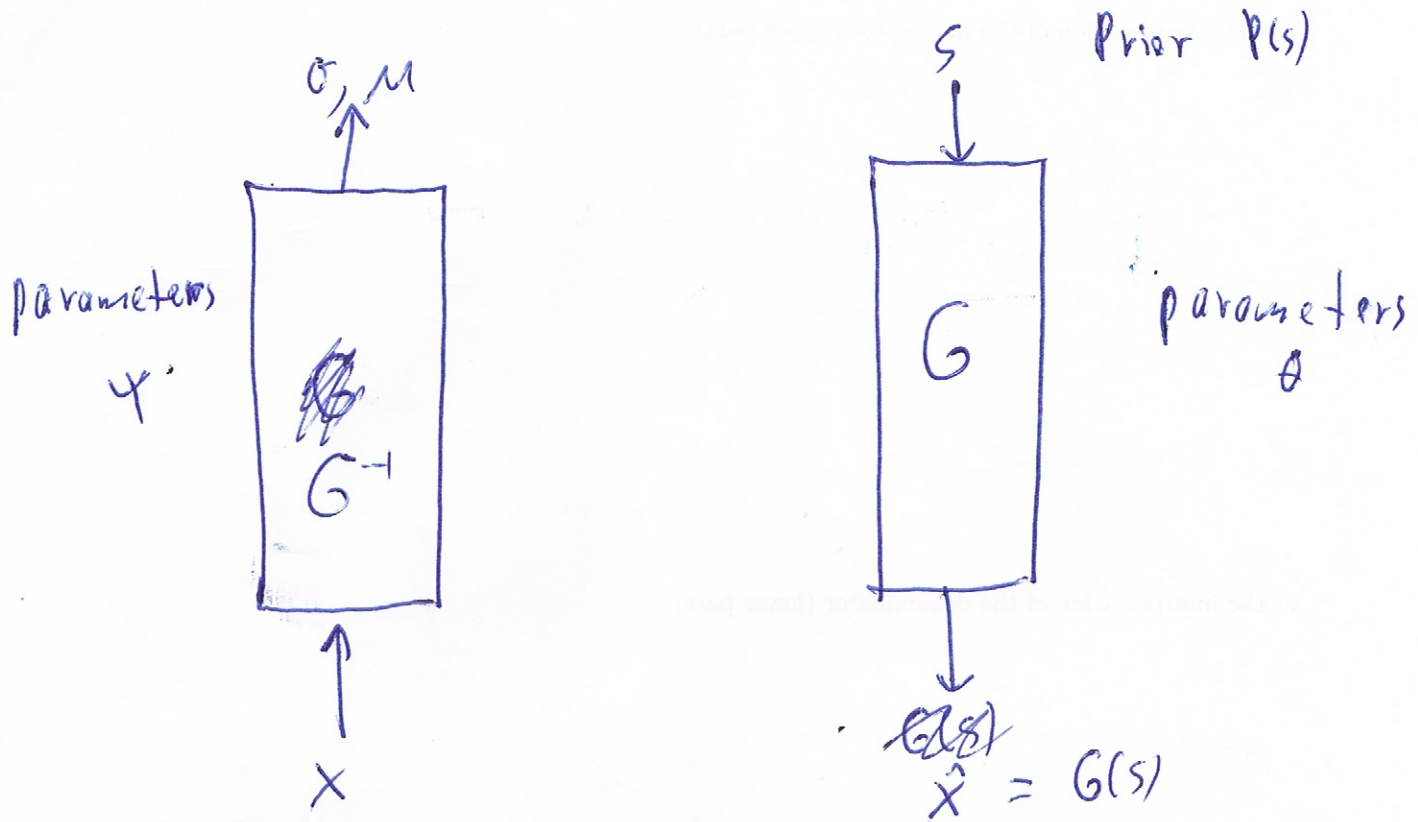
$$\approx \frac{1}{L} \sum_{\ell=1}^L f'(\mu + \sigma \xi_{\ell}) (\nabla_{\Psi} \mu + \xi_{\ell} \nabla_{\Psi} \sigma)$$

$$\text{where } \xi_{\ell} \sim \mathcal{N}(0,1)$$

Note: f must be differentiable

Variational Auto encoders

(13)



$$p(x|\hat{x}) = \mathcal{N}(x|\hat{x}, \sigma^2)$$

typically $p(z) = \mathcal{N}(z|0, 1)$

$$\begin{aligned} \log p(x|z) &\approx \log p(x|\hat{x}) \\ &= -\frac{(x - \hat{x})^2}{2\sigma^2} + C \\ &= -\frac{(x - G(z))^2}{2\sigma^2} + C \end{aligned}$$

$$\mathcal{L} = \sum_n \mathbb{E}_{s \sim q(s|x^n)} \log \frac{p_\theta(x^n, s)}{q_\psi(s|x^n)}$$

$$= \sum_n \mathbb{E}_{s \sim q(s|x^n)} \log p_\theta(x^n, s)$$

$$- \sum_n \mathbb{E}_{s \sim q(s|x^n)} \log q_\psi(s|x^n)$$

$$= \sum_n \mathbb{E}_{s \sim q(s|x^n)} \log p_\theta(x^n | s)$$

$$+ \sum_n H(q(s|x^n)) \quad \leftarrow \text{Entropy of } q(s|x^n) \text{ (usually a simple formula)}$$

eg. if $q(s|x) = \mathcal{N}(s|\mu, \sigma)$

then $H(s|x) = \log \sigma + C$

where $C = \frac{1}{2} + \frac{1}{2} \log(2\pi)$

$$\begin{aligned} \therefore \nabla_{\psi} \mathcal{L} &= \sum_n \nabla_{\psi} \mathbb{E}_{s \sim q(s|x^n)} \overbrace{\log p_{\theta}(x^n; s)}^{f_n(s)} \\ &\quad + \sum_n \nabla_{\psi} H(q(s|x^n)) \end{aligned}$$

$$= \sum_n \nabla_{\psi} \mathbb{E}_{s \sim q(s|x^n)} f_n(s)$$

$$+ \sum_n \nabla_{\psi} \log \sigma_n$$

where $s|x^n \sim \mathcal{N}(s|\mu_n, \sigma_n)$

$$= \sum_n \nabla_{\psi} \mathbb{E}_{\epsilon \sim \mathcal{N}(0,1)} f_n(\mu_n + \sigma_n \epsilon) \} \text{ hard.}$$

$$+ \sum_n \nabla_{\psi} \log \sigma_n \} \begin{array}{l} \text{easy} \\ \text{(back prop through } G^{-1}) \end{array}$$

The hard part

(15)

$$\sum_n \nabla_{\Psi} E_{\xi \sim \mathcal{N}(0,1)} f_n(\mu_n + \sigma_n \xi)$$

$$= \sum_n E_{\xi \sim \mathcal{N}(0,1)} \nabla_{\Psi} f_n(\mu_n + \sigma_n \xi)$$

$$= \sum_n E_{\xi \sim \mathcal{N}(0,1)} f'_n(\mu_n + \sigma_n \xi) \cdot \nabla_{\Psi} (\mu_n + \sigma_n \xi)$$

$$\approx \sum_n \frac{1}{L} \sum_{\ell=1}^L f'_n(\mu_n + \sigma_n \xi_n^{\ell}) \nabla_{\Psi} (\mu_n + \sigma_n \xi_n^{\ell})$$

where $\xi_n^{\ell} \sim \mathcal{N}(0,1)$

(Can use $L=1$ when N is large)

$$\approx \sum_n f'_n(\mu_n + \sigma_n \xi_n) (\nabla_{\Psi} \mu_n + \xi_n \nabla_{\Psi} \sigma_n)$$

↑
backprop
through G

↑
backprop
through G^{-1}

where $\xi_n \sim \mathcal{N}(0,1)$

Discrete Case

(s is discrete)

Recall the general problem:

evaluate $\nabla_{\Psi} E_{s \sim q} f(s)$

where (the parameters of) q depends on Ψ .

Simple example:

$$s \in \{0, 1\}$$

$$q(s=1) = \alpha$$

$$q(s=0) = 1 - \alpha$$

where α depends on Ψ .

eg, α is the output of a neural net with weights Ψ .

(17)

Note that even though s is discrete, the expectation $\mathbb{E}$ over s is differentiable as long as α is differentiable in ψ .

eg. in the simple example above,

$$\begin{aligned}
 \mathbb{E}_{s \sim q} f(s) &= \cancel{q(1)f(1)} + \sum_s q(s) f(s) \\
 &= q(1)f(1) + q(0)f(0) \\
 &= \alpha \cdot f(1) + (1-\alpha) \cdot f(0) \\
 &= \alpha \cdot [f(1) - f(0)] + f(0)
 \end{aligned}$$

$$\therefore \nabla_{\psi} \mathbb{E}_{s \sim q} f(s) = [f(1) - f(0)] \cdot \nabla_{\psi} \alpha$$

The problem is when s is a vector.
Then the expectation (a sum) has a
huge number of terms (eg, 2^{100}), far too
many to enumerate.

But, because it is an expectation, we
can try to approximate it using sampling.

~~Alas, for discrete distributions, this
almost always returns 0 as the estimate of
the gradient.~~

$$\nabla_{\psi} E_{s \sim q} f(s) \approx \nabla_{\psi} \frac{1}{L} \sum_{l=1}^L f(s_l) \quad \text{where } s_l \sim q$$

Since q depends on ψ , how do we
differentiate through the sampling process?

re Parameterization Trick doesn't work

example: $s \in \{0, 1\}$

$$q(s=1) = \alpha$$

$$q(s=0) = 1 - \alpha$$

} α depends on ψ

To use reparam. trick, must make s a function of α & some other random variable, ϵ , ind. of ψ

eg. $s = [\epsilon < \alpha]$ where $\epsilon \sim U(0, 1)$

$$\therefore \Pr(\epsilon < \alpha) = \alpha$$

note $[\text{True}] = 1$

$[\text{False}] = 0$

$$\therefore \nabla_{\alpha} E_{\xi \sim q} f(\xi)$$

$$= \nabla_{\alpha} E_{\xi \sim u(0,1)} f([\xi \leq \alpha])$$

$$= E_{\xi \sim u(0,1)} \nabla_{\alpha} f([\xi \leq \alpha])$$

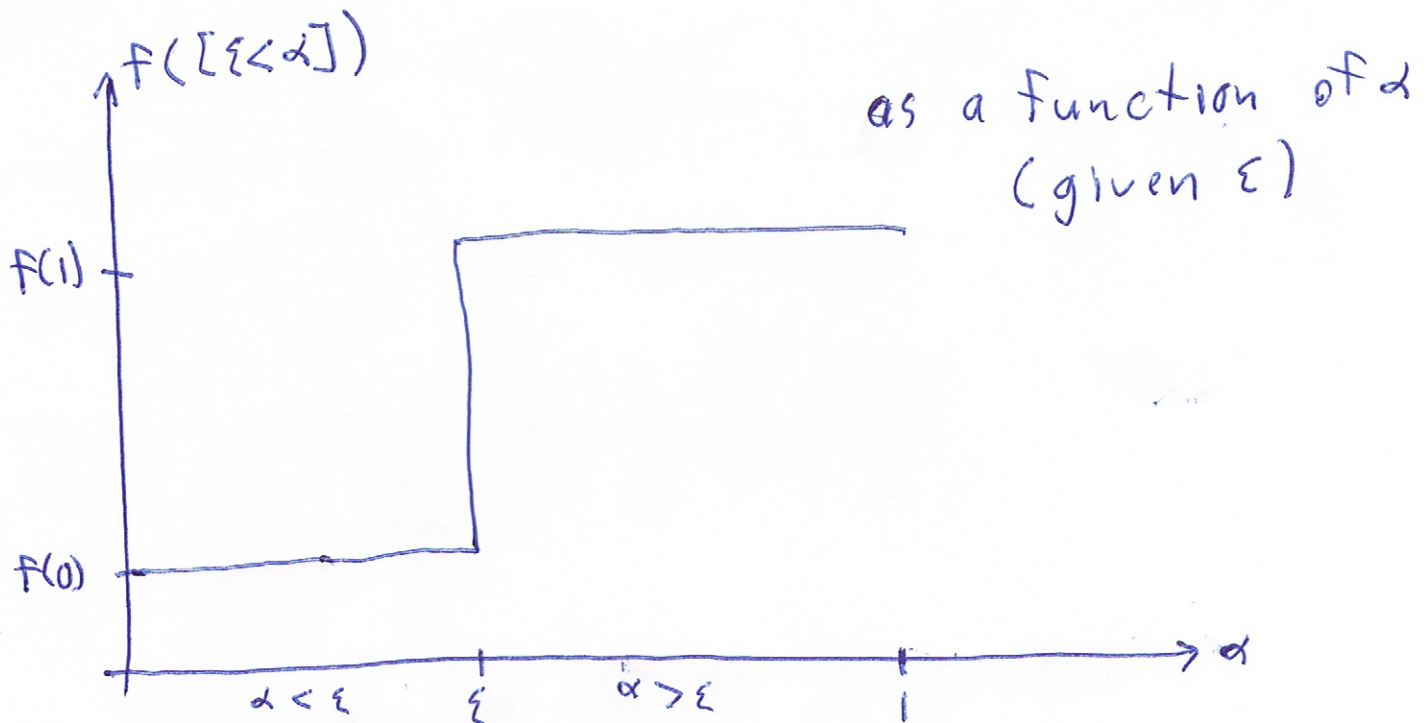
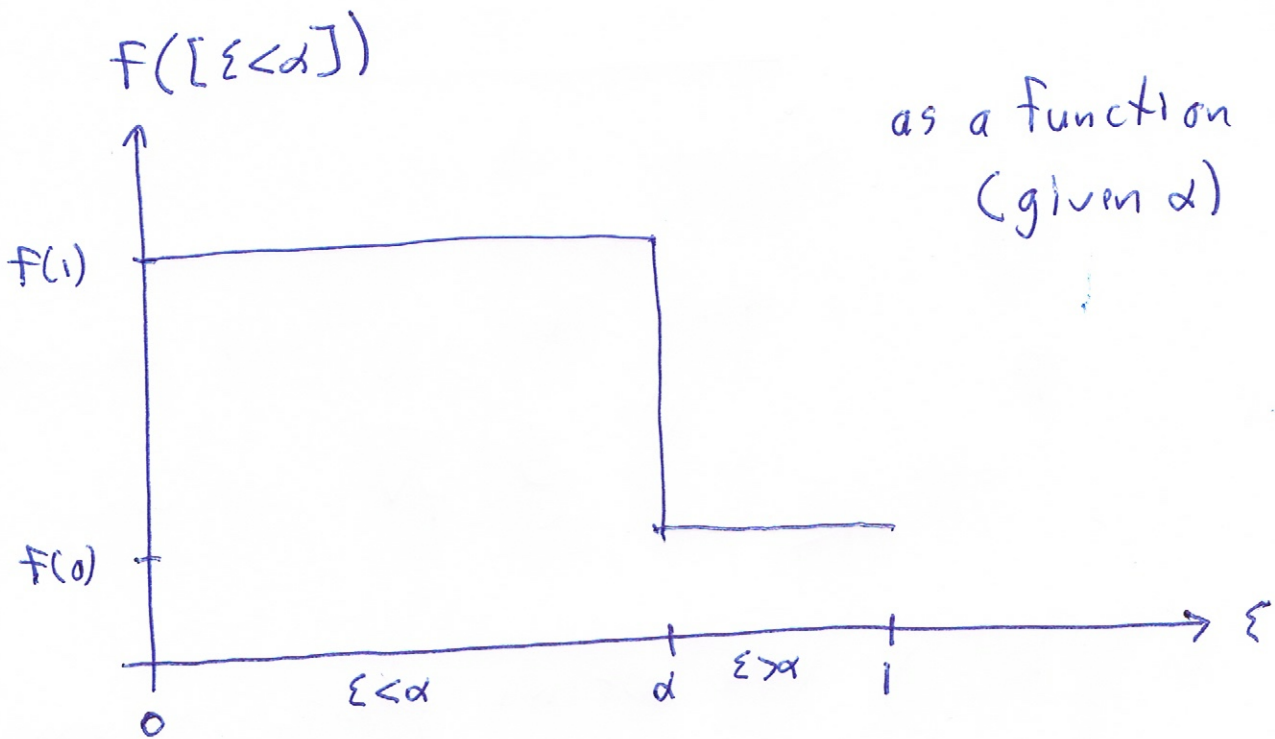
$$= E_{\xi \sim u(0,1)} \frac{\partial}{\partial \alpha} f([\xi \leq \alpha]) \cdot \frac{\partial \alpha}{\partial \alpha}$$

at try
param.
trick

but $f([\xi \leq \alpha])$ is constant a.e.

$$f([\xi \leq \alpha]) = \begin{cases} f(1) & \text{if } \xi < \alpha \\ f(0), & \text{o.w.} \end{cases}$$

$$\therefore \frac{\partial}{\partial \alpha} f([\xi \leq \alpha]) = \begin{cases} 0 & \text{if } \xi \neq \alpha \\ \infty & \text{a.e. if } \xi = \alpha \end{cases}$$



$$\therefore \frac{\partial F([\varepsilon < \alpha])}{\partial \alpha} = 0 \text{ almost everywhere}$$

$$\therefore \nabla_{\psi} E_{\text{single}} f(s)$$

$$= E_{\xi \sim u(0,1)} \frac{\partial f([\xi \leq \alpha])}{\partial \alpha} \cdot \frac{\partial \alpha}{\partial \psi}$$

$$\approx \frac{1}{L} \sum_{\ell=1}^L \frac{\partial f([\xi_{\ell} \leq \alpha])}{\partial \alpha} \frac{\partial \alpha}{\partial \psi}$$

where $\xi_{\ell} \sim u(0,1)$

$$= \begin{cases} 0 & \text{if all } \xi_{\ell} \neq \alpha \\ \infty & \text{if some } \xi_{\ell} = \alpha \end{cases}$$

ie, a very high variance estimate.

what can we do?