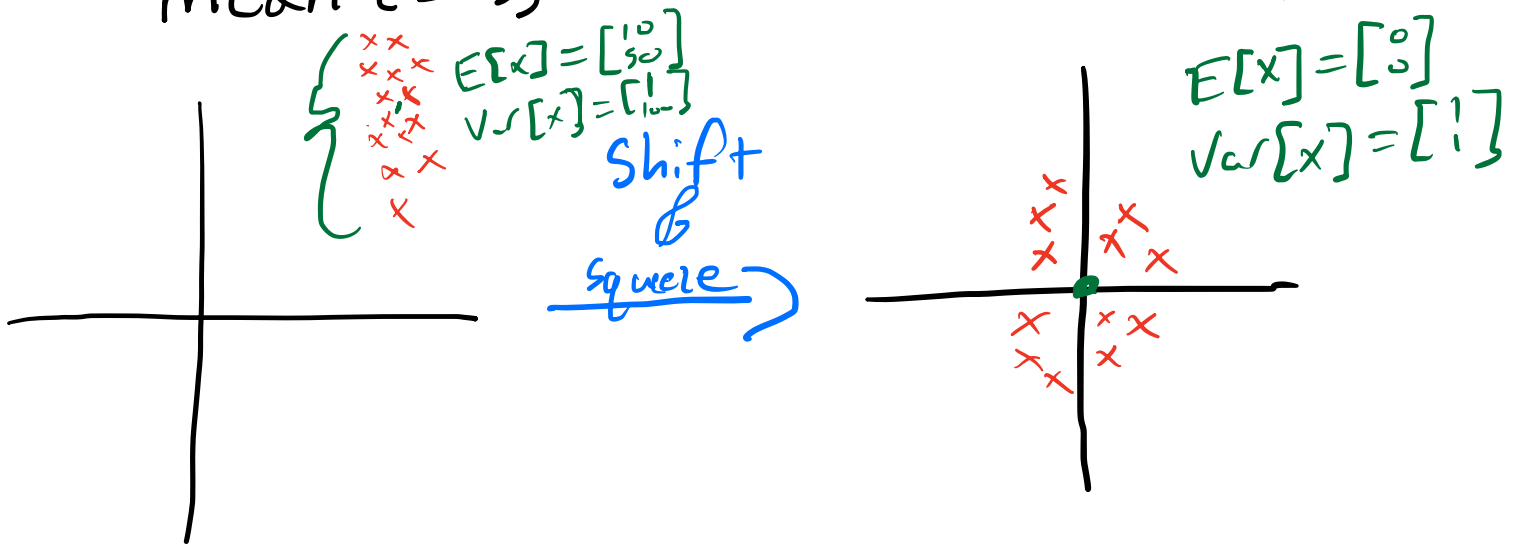


Last time

Normalization: keep data at consistent mean ($E[x]$) and Variance ($\text{Var}[x]$).



Motivation

1) Gradient scale

$$\left| \frac{dL}{dw_1} \right| = |x| |\sigma'(xw_1 + b_1)| \prod_{i=2}^L \left| \frac{d\phi_i}{d\phi_{i-1}} \right| \dots$$

(Prevent scale of data from affecting gradient)

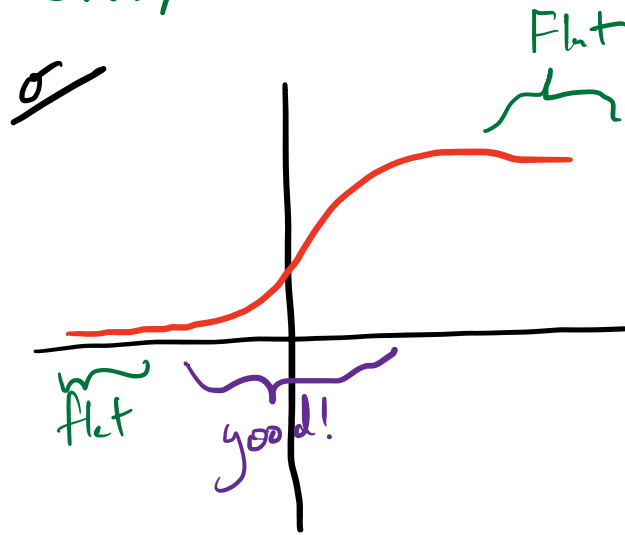
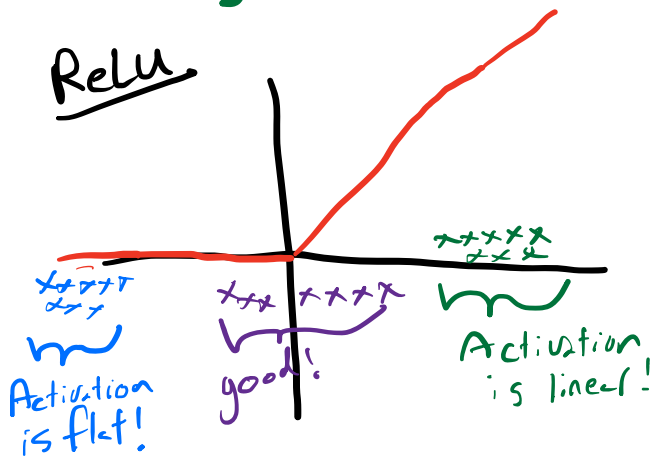
2) Scale mismatch

e.g. $x_1 = \begin{bmatrix} 2152 \\ 0.02 \end{bmatrix}$ $x_2 = \begin{bmatrix} 1.92 \\ 0.017 \end{bmatrix} \dots$

We saw why this is a problem!

Motivation

3) Activations only make sense need 0



Normalization

$$\text{Norm}(x) = \frac{x - E[x]}{\sqrt{\text{Var}[x] + \epsilon}}$$

Prevent div. by 0!
 $\epsilon \ll 1$

Estimate $E[x]$, $\text{Var}[x]$ from dataset or Batch

$$E[x] \approx \bar{x} = \frac{1}{N} \sum_{i=1}^N x_i$$

$$\text{Var}[x] \approx s^2 = \frac{1}{N-1} \sum_{i=1}^N (x_i - \bar{x})^2$$

Unbiased estimator

$$E[\text{Norm}(x)] = E\left[\frac{x - E[x]}{\sqrt{\text{Var}[x] + \epsilon}}\right]$$
$$= \frac{1}{\sqrt{\text{Var}[x] + \epsilon}} [E[x] - E[x]] = 0$$

Batch

Normalization

$$\text{Batch Norm}(x) = \frac{x - E[x]}{\sqrt{\text{Var}[x] + \epsilon}} \quad \text{Prevent div. by 0!}$$

Estimate $E[x]$, $\text{Var}[x]$ from dataset of Batch

← Batch sampled for SGD

$$E[x] \approx \bar{x} = \frac{1}{B} \sum_{i=1}^B x_i$$

$$\text{Var}[x] \approx s^2 = \frac{1}{B-1} \sum_{i=1}^B (x_i - \bar{x})^2$$

Unbiased estimator

Deep Network

$$f(x) = \phi_l(\phi_{l-1}(\dots \phi_1(x)))w_0 + b_0$$

Composing feature functions

e.g.

$$f(x) = \sigma(\sigma(\dots \sigma(x^T w_1 + b_1)^T w_2 + b_2) \dots)^T w_l + b_l$$

or

$$\phi_1 = \sigma(x^T w_1 + b_1)$$

$$\phi_2 = \sigma(\phi_1^T w_2 + b_2)$$

$\vdots$

$$\phi_l = \sigma(\phi_{l-1}^T w_l + b_l)$$

$$f = \phi_l^T w_0 + b_0$$

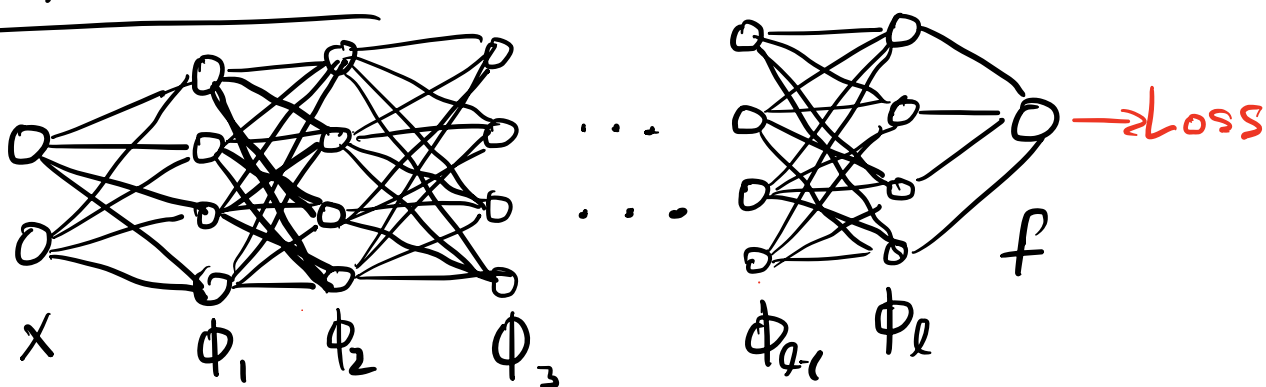
$$L = \text{Loss}(f, y)$$

e.g. $\downarrow$
 $(f - y)^2$ MSE

↳ Is our internal representation normalized? Consistent?

Deep Network

l -layer network



or

$$\phi_1 = \sigma(x^T \underline{w_1} + b_1)$$

$$\phi_2 = \sigma(\phi_1^T \underline{w_2} + b_2)$$

$\vdots$

$$\phi_l = \sigma(\underline{\phi_{l-1}}^T w_l + b_l)$$

$$f = \phi_l^T w_o + b_o$$

$$L = \text{Loss}(f, y)$$

$\downarrow$
e.g. $(f - y)^2$

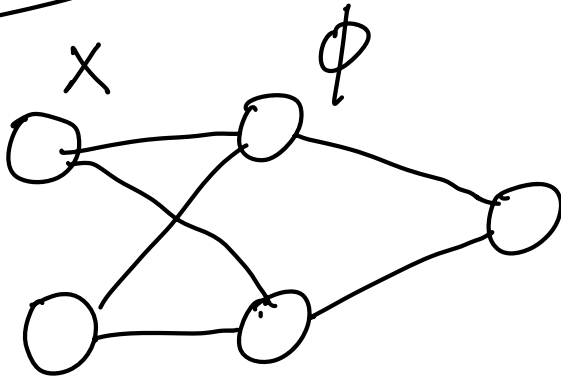
$$E[x] = 0 \quad \text{Var}[x] = 1 \quad \text{w/ Norm.}$$

$$E[\phi_{l-1}] = ? \quad \text{Var}[\phi_{l-1}] = ?$$

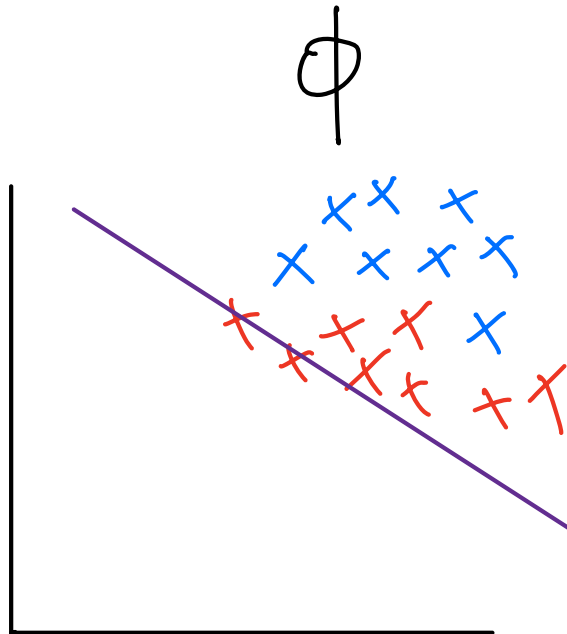
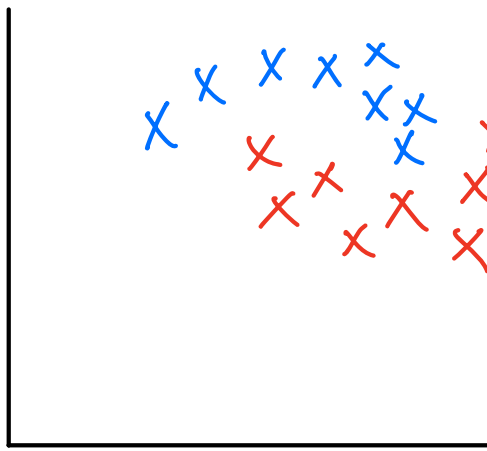
Depend on $w_1 \dots w_{l-1}$ which could change

Internal covariate shift

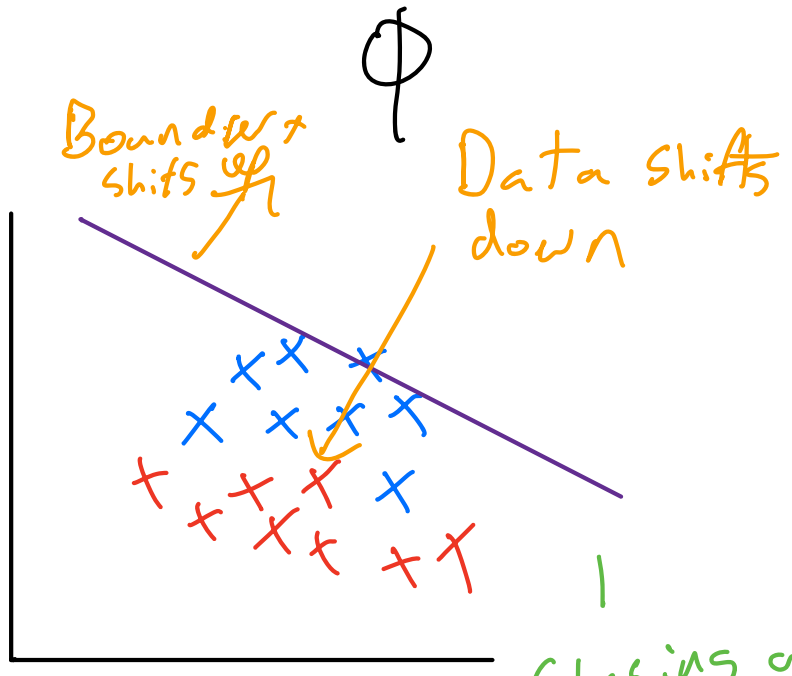
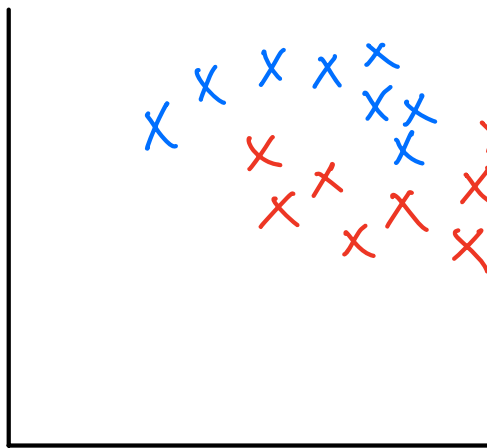
Simple network



Step k

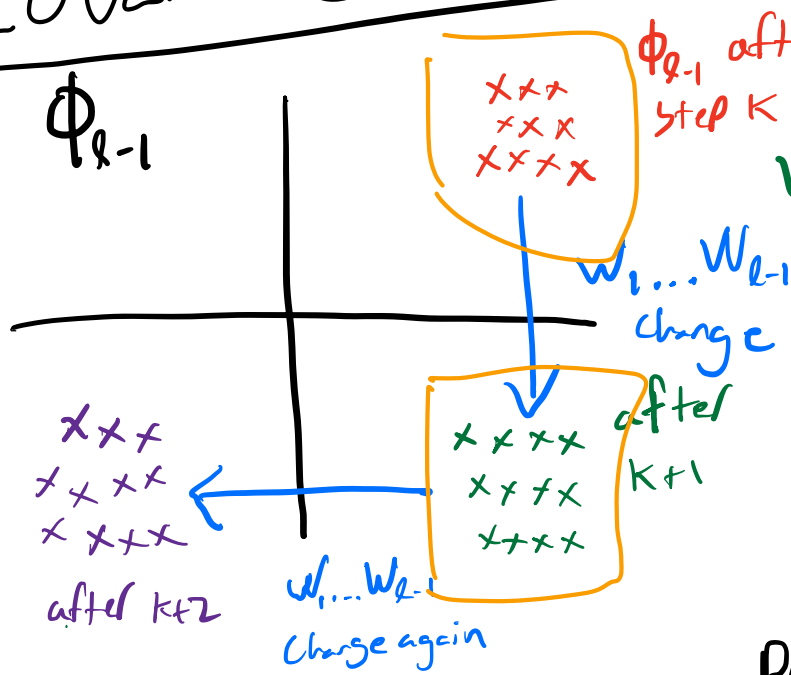


Step k+1



Chasing or moving target!

Covariate shift



Layer l

$$\phi_l = \sigma(\phi_{l-1} w_l + b_l)$$

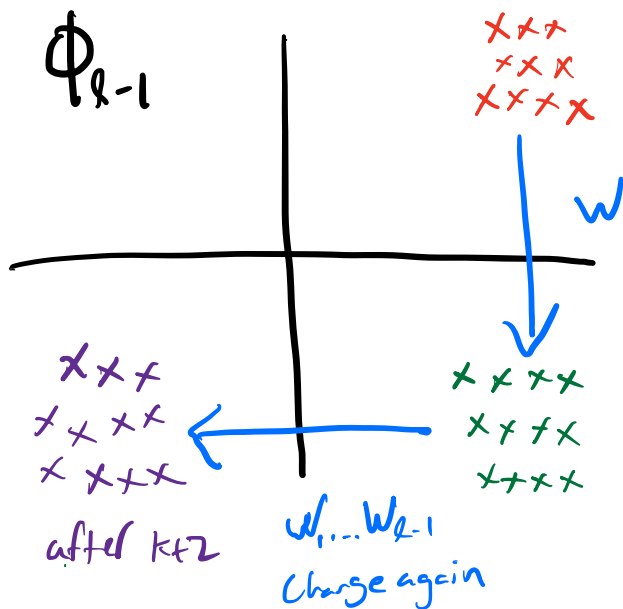
$$w_l^{(k+1)} \leftarrow w_l^{(k)} - \alpha \phi_{l-1}^{(k)} \sigma(\phi_{l-1}^{(k)} w_l^{(k)} + b_l)$$

update using old $\phi_{l-1}^{(k)}$

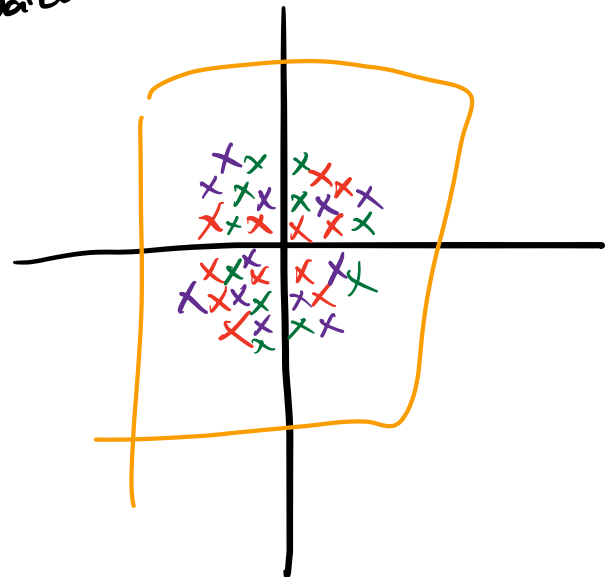
$$f(x) = \sigma(\phi_{l-1}^{(k+1)} w_l^{(k+1)} + b_l) \dots$$

Predict using new $\phi_{l-1}^{(k+1)}$

Batch Norm

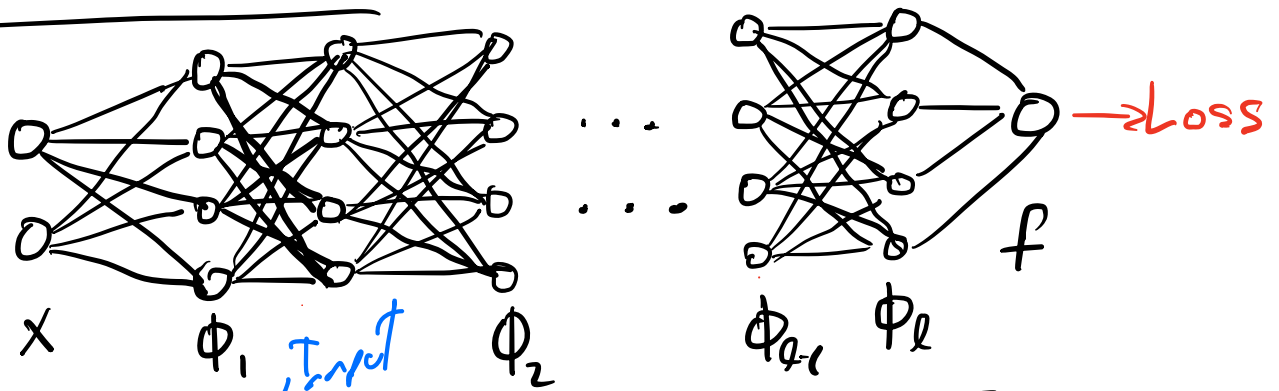


BatchNorm(ϕ_{l-1})



Can apply Batch Norm IN the network

Deep Network w/ Batch Norm!



$$\begin{aligned} \phi_1 &= \sigma(\text{BN}(x)^T w_1 + b_1) & f &= \text{BN}(\phi_l)^T w_o + b_o \\ \phi_2 &= \sigma(\text{BN}(\phi_1)^T w_2 + b_2) & L &= \text{Loss}(f, y) \\ &\uparrow \text{output of } \phi_{l-1} & &\downarrow \\ & & & \text{e.g. } (f-y)^2 \end{aligned}$$

$$E[x] = 0 \quad \text{Var}[x] = 1 \quad \text{w/ Norm.}$$

$$E[\text{BN}(\phi_{l-1})] = 0 \quad \text{Var}[\text{BN}(\phi_{l-1})] = 1$$

Deep Network w/ Batch Norm

$$E[\text{BN}(\phi_{l-1})] = 0 \quad \text{Var}[\text{BN}(\phi_{l-1})] = 1$$

$E[\phi_{l-1}]$, $\text{Var}[\phi_{l-1}]$ still change every step

$\therefore$ We need to use **Batch Norm** and
update $\bar{x} \approx E[x]$ and $s^2 \approx \text{Var}[x]$ at every
step
($E[\phi_{l-1}]$ for layer l)

Test time

Problem: At test time we might want
to make a prediction on a single obs.

Single obs. $\rightarrow B=1$

$$\bar{x} = \frac{1}{1} x_1$$

$$x_1 - \bar{x} = 0!$$

$$s^2 = \frac{1}{1-1} (x_1 - \bar{x})^2$$

$\underbrace{\quad}_{\text{divide by 0!}}$
 $= 0 \rightarrow \text{divide by 0!}$

Batch Norm at test time

In practice: Maintain Running Average of mean and Variance estimates as we go. (Using EMA!)

At step k : $\bar{X}_{Avg.}^{(k)} \leftarrow \beta \bar{X}_{Avg.}^{(k-1)} + (1-\beta) \bar{X}^{(k)}$ $\beta \in [0, 1]$
typical 0.9

$$S_{Avg.}^2 \leftarrow \beta S_{Avg.}^2 + (1-\beta) S^{(k)2}$$

Left using current batch

At test:

$$\text{Batch Norm}_{\text{test}}(x) = \frac{x - \bar{X}_{Avg.}}{\sqrt{S_{Avg.}^2 + \epsilon}}$$

Use the same estimates every time we evaluate

Layer Norm

What if we don't want to do separate things at train and test time?

Batch norm

$$\begin{bmatrix} x_{11} \\ x_{12} \\ x_{13} \end{bmatrix}, \begin{bmatrix} x_{21} \\ x_{22} \\ x_{23} \end{bmatrix}, \begin{bmatrix} x_{31} \\ x_{32} \\ x_{33} \end{bmatrix}, \dots$$

Find $\bar{x} \approx E[x]$, $s \approx \text{Var}[x]$ by Averaging over observations for each dimension

Layer norm

$$\begin{bmatrix} x_{11} \\ x_{12} \\ x_{13} \end{bmatrix}, \begin{bmatrix} x_{21} \\ x_{22} \\ x_{23} \end{bmatrix}, \begin{bmatrix} x_{31} \\ x_{32} \\ x_{33} \end{bmatrix}, \dots$$

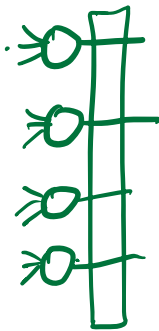
$E[x_{1*}] = 0$ $\text{Var}[x_{1*}] = 1$

Alt. Average over dimensions within an observation

Layer Norm

$$\text{Layer Norm}(x_i) = \frac{x_i - \bar{x}_i}{\sqrt{s^2 + \epsilon}}$$
$$\bar{x}_i = \frac{1}{d} \sum_{j=1}^d x_{ij} \quad s^2 = \frac{1}{d-1} \sum_{j=1}^d (x_{ij} - \bar{x}_i)^2$$

In network



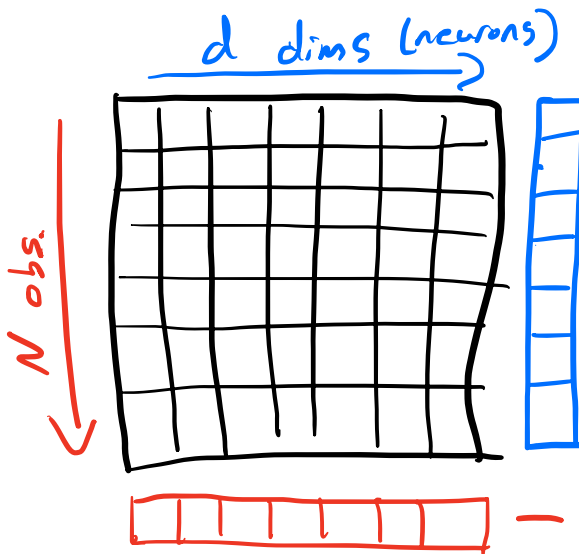
$$E[x] = 0, \text{var}[x] = 1$$

Can be applied w/
Batch size 1! As long
as $d \geq 1$

Batch Norm vs. Layer Norm

Input matrix

$\underline{X}$ $\rightarrow$
 $N \times d$



Layer norm
statistics per row

Batch norm
statistics per column

Batch Norm Usually preferred, but not always possible
[preserves more info about each observation]

Batch Norm caveats

Covariate shift explanation is controversial!

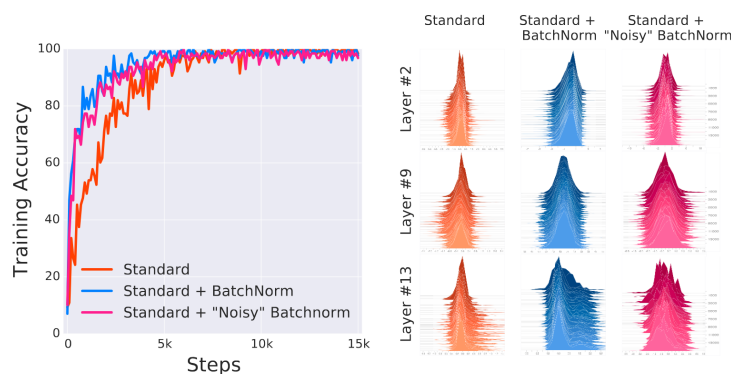


Figure 2: Connections between distributional stability and BatchNorm performance: We compare VGG networks trained without BatchNorm (Standard), with BatchNorm (Standard + BatchNorm) and with explicit “covariate shift” added to BatchNorm layers (Standard + “Noisy” BatchNorm). In the later case, we induce distributional instability by adding *time-varying, non-zero* mean and *non-unit* variance noise independently to each batch normalized activation. The “noisy” BatchNorm model nearly matches the performance of standard BatchNorm model, despite complete distributional instability. We sampled activations of a given layer and visualized their distributions (also cf. Figure 7).

Alternative claim: smooths Loss landscape

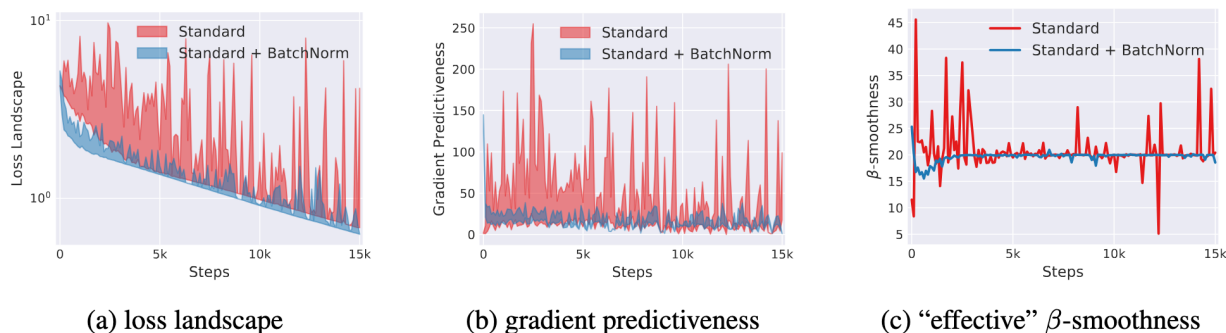


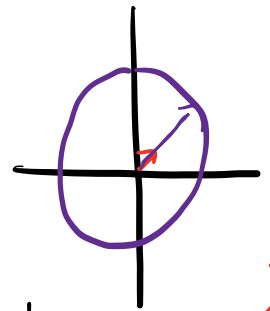
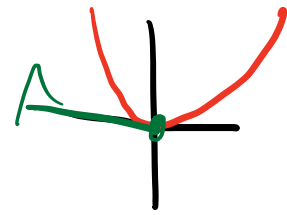
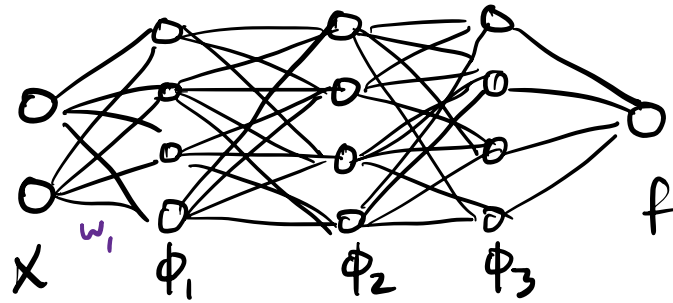
Figure 4: Analysis of the optimization landscape of VGG networks. At a particular training step, we measure the variation (shaded region) in loss (a) and ℓ_2 changes in the gradient (b) as we move in the gradient direction. The “effective” β -smoothness (c) refers to the maximum difference (in ℓ_2 -norm) in gradient over distance moved in that direction. There is a clear improvement in all of these measures in networks with BatchNorm, indicating a more well-behaved loss landscape. (Here, we cap the maximum distance to be $\eta = 0.4 \times$ the gradient since for larger steps the standard network just performs worse (see Figure 1). BatchNorm however continues to provide smoothing for even larger distances.) Note that these results are supported by our theoretical findings (Section 4).

Some debate over ordering

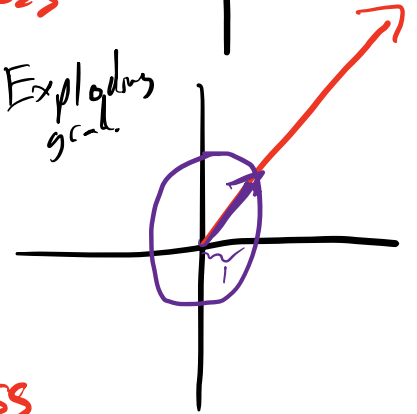
$$\sigma(\text{BN}(x)w + b) \text{ vs. } \sigma(\text{BN}(wx + b))$$

↖ More common these days

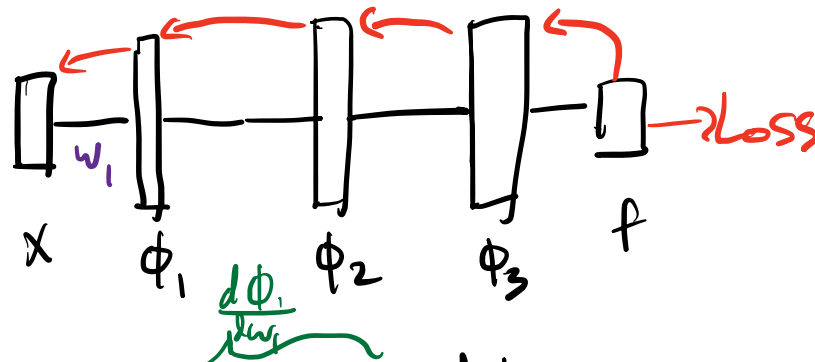
Residual Networks



Exploding grad.



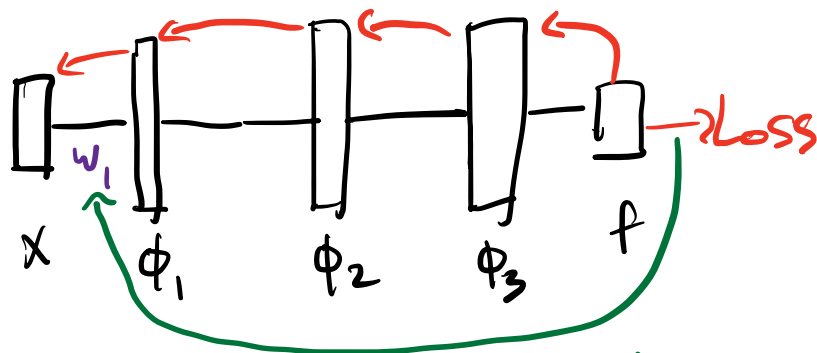
Simplified



$$\frac{dL}{dw_1} = x \sigma'(x^T w_1) \frac{d\phi_2}{d\phi_1} \cdot \frac{d\phi_3}{d\phi_2} \cdots \frac{df}{d\phi_3} \cdot \frac{dLoss}{df}$$

Info from Loss gets diluted
(vanishing gradients)

Residual Networks



What if we had a direct connection?

$$\frac{dL}{dw_1} = x \sigma'(x^T w_1) \underbrace{\frac{d\phi_2}{d\phi_1} \cdot \frac{d\phi_3}{d\phi_2} \cdots \frac{df}{d\phi_3}}_{\left| \frac{d\phi_i}{d\phi_{i-1}} \right| < 1 \rightarrow 0} \cdot \frac{d\text{Loss}}{df} + \underbrace{\frac{dL}{d\phi_1} \cdot \frac{d\phi_1}{dw_1}}_{\text{Loss depends directly on } \phi_1}$$

Residual function

$$\hat{f}(x) \rightarrow f(x) = \hat{f}(x) + x$$

Add input back to output!

Residual Layer

$$\text{Layer: } \phi_i = \sigma(\underbrace{\phi_{i-1}^T w_i + b_i}_{\text{next value}})$$

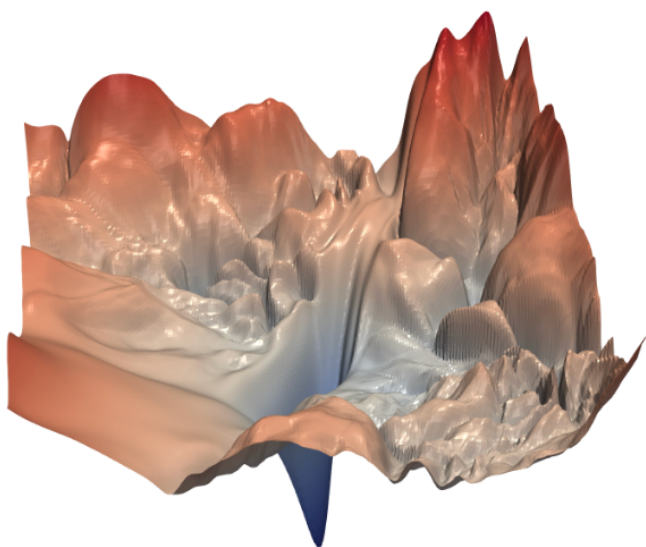
$$\text{Residual Layer: } \phi_i = \sigma(\phi_{i-1}^T w_i + b_i) + \phi_{i-1}$$

$$\frac{d\phi_i}{d\phi_{i-1}} = \frac{d}{d\phi_{i-1}} \left[\sigma(\phi_{i-1}^T w_i + b_i) + \phi_{i-1} \right]$$

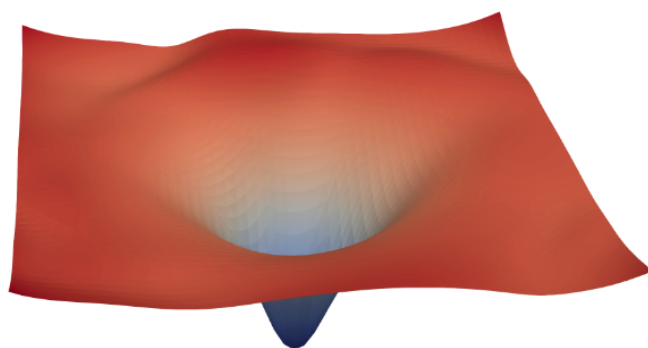
$$= \underbrace{\sigma'(\phi_{i-1}^T w_i + b_i) w_i}_{} + 1$$

$$\frac{d\phi_i}{d\phi_{i-1}} = \frac{d\hat{\phi}_i}{d\phi_{i-1}} + 1$$

$$\begin{aligned}
 \rightarrow \frac{dL}{dw_1} &= \frac{d\phi_1}{dw_1} \underbrace{\left(\frac{d\hat{\phi}_2}{d\phi_1} + 1 \right) \left(\frac{d\hat{\phi}_3}{d\phi_2} + 1 \right) \dots}_{\prod_{i=2}^L \left(\frac{d\hat{\phi}_i}{d\phi_{i-1}} + 1 \right)} \cdot \underbrace{\frac{df}{d\phi_L}}_{\frac{dL_{\text{loss}}}{df}} \frac{dL_{\text{loss}}}{df} \\
 &\quad \underbrace{\left(\frac{d\hat{\phi}_2}{d\phi_1} + 1 \right) \left(\frac{d\hat{\phi}_3}{d\phi_2} + 1 \right)}_{+1}
 \end{aligned}$$



(a) without skip connections



(b) with skip connections

Figure 1: The loss surfaces of ResNet-56 with/without skip connections. The proposed filter normalization scheme is used to enable comparisons of sharpness/flatness between the two figures.

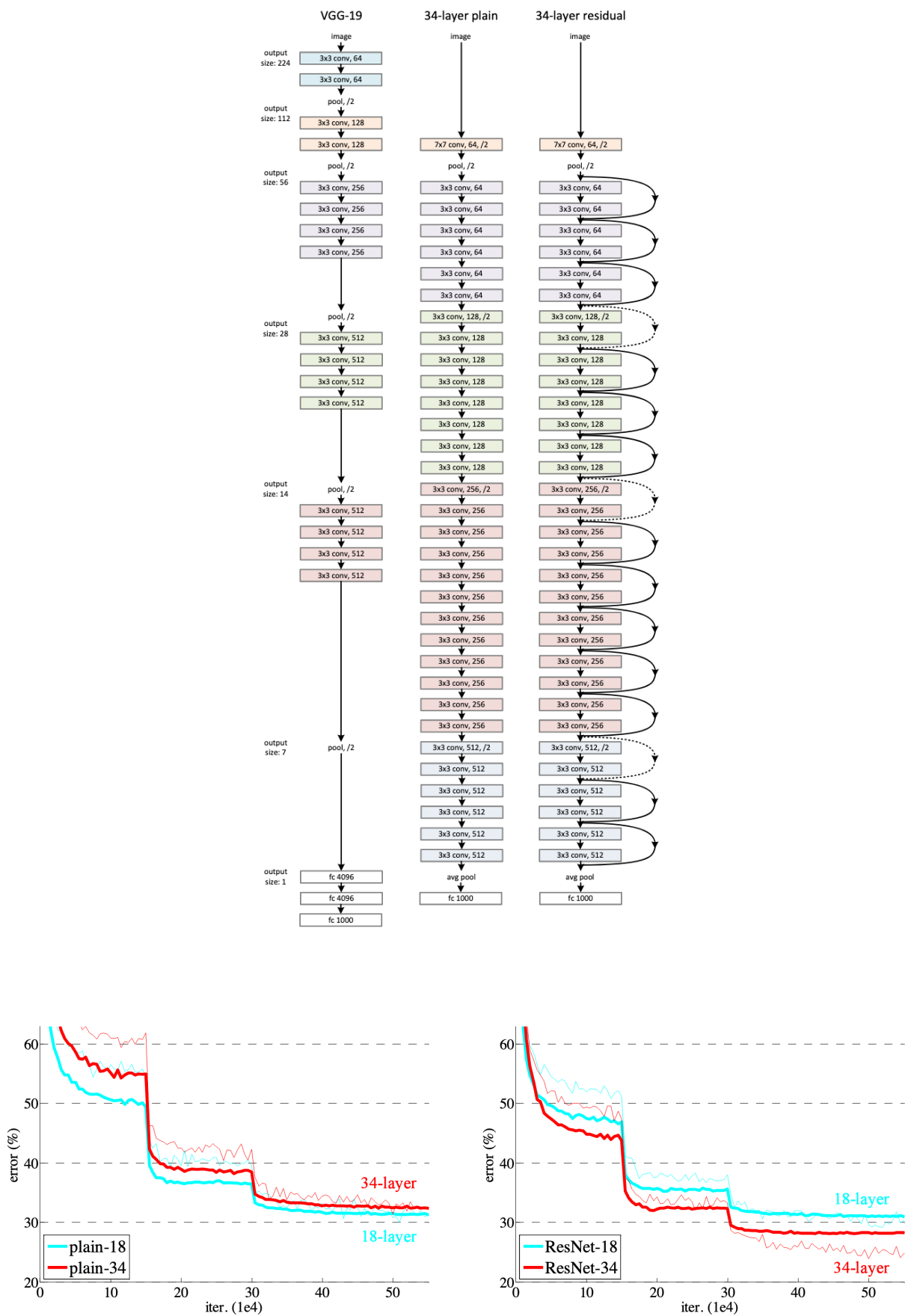


Figure 4. Training on **ImageNet**. Thin curves denote training error, and bold curves denote validation error of the center crops. Left: plain networks of 18 and 34 layers. Right: ResNets of 18 and 34 layers. In this plot, the residual networks have no extra parameter compared to their plain counterparts.