

Initialization

Gabriel Hope

Adam

Can we combine adaptive scaling and momentum?

Adam

Update *velocity*

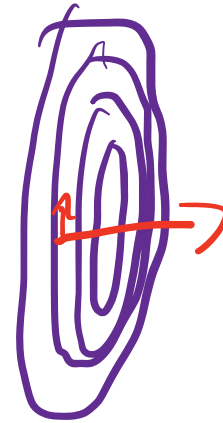
$$\mathbf{v}^{(k+1)} \leftarrow \beta_1 \mathbf{v}^{(k)} + (1 - \beta_1) \nabla_{\mathbf{w}} \mathbf{Loss}(\mathbf{w}^{(k)}, \mathbf{X}, \mathbf{y})$$

Update *scaling*

$$\mathbf{s}^{(k+1)} \leftarrow \beta_2 \mathbf{s}^{(k)} + (1 - \beta_2) (\nabla_{\mathbf{w}} \mathbf{Loss}(\mathbf{w}^{(k)}, \mathbf{X}, \mathbf{y}))^2$$

Update *weights*

$$\mathbf{w}^{(k+1)} \leftarrow \mathbf{w}^{(k)} - \alpha \frac{\mathbf{v}^{(k+1)}}{\sqrt{\mathbf{s}^{(k+1)} + \epsilon}}$$



Adam

Update *velocity*

$$\mathbf{v}^{(k+1)} \leftarrow \beta_1 \mathbf{v}^{(k)} + (1 - \beta_1) \nabla_{\mathbf{w}} \text{Loss}(\mathbf{w}^{(k)}, \mathbf{X}, \mathbf{y})$$

Update *scaling*

$$\mathbf{s}^{(k+1)} \leftarrow \beta_2 \mathbf{s}^{(k)} + (1 - \beta_2) (\nabla_{\mathbf{w}} \text{Loss}(\mathbf{w}^{(k)}, \mathbf{X}, \mathbf{y}))^2$$

Modified weight update:

$$\mathbf{w}^{(k+1)} \leftarrow \mathbf{w}^{(k)} - \boxed{\alpha} \frac{\mathbf{v}^{(k+1)}}{\sqrt{\frac{\mathbf{s}^{(k+1)}}{(1-\beta_2^k)} + \epsilon}}$$

learning rate
velocity
scale

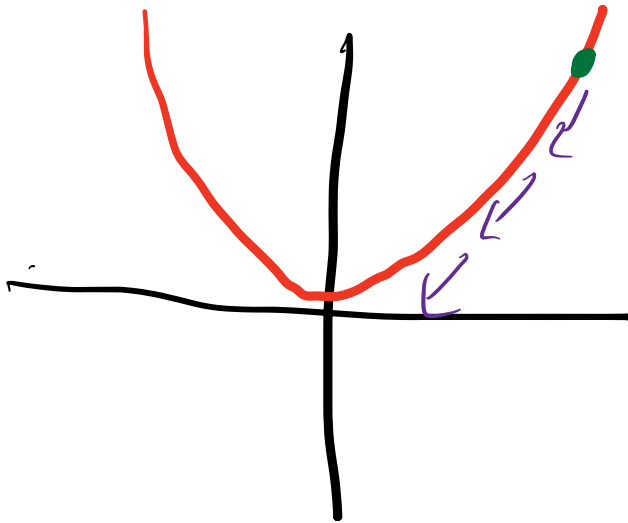
Adam

At step 0:

$$v^{(k+1)} = \beta_1 \frac{v^{(k)}}{0} + (1 - \beta_1) \nabla$$

$k=0$
 $\underline{v^{(0)} = 0}, \quad s^{(0)} = 0$

$$\frac{v^{(k+1)}}{(1 - \beta_1^k)} = \frac{\beta_1 \mathbf{0} + (1 - \beta_1) \nabla_{\mathbf{w}} \text{Loss}(\mathbf{w}^{(k)}, \mathbf{X}, \mathbf{y})}{(1 - \beta_1^k)} = \underline{\nabla_{\mathbf{w}} \text{Loss}(\mathbf{w}^{(k)}, \mathbf{X}, \mathbf{y})}$$



as $k \rightarrow \infty$

$$(1 - \beta^k) \rightarrow 1$$

e.g.

$$\beta = 0.9$$

$$[0.9]^{100} \approx 0$$

$$\rightarrow [1 - \beta^{100}] \approx 1$$

Summary of gradient descent issues

Updates are too slow

- Stochastic/minibatch gradient descent

SGD gradients are very noisy (high variance)

- Increase batch size, use momentum

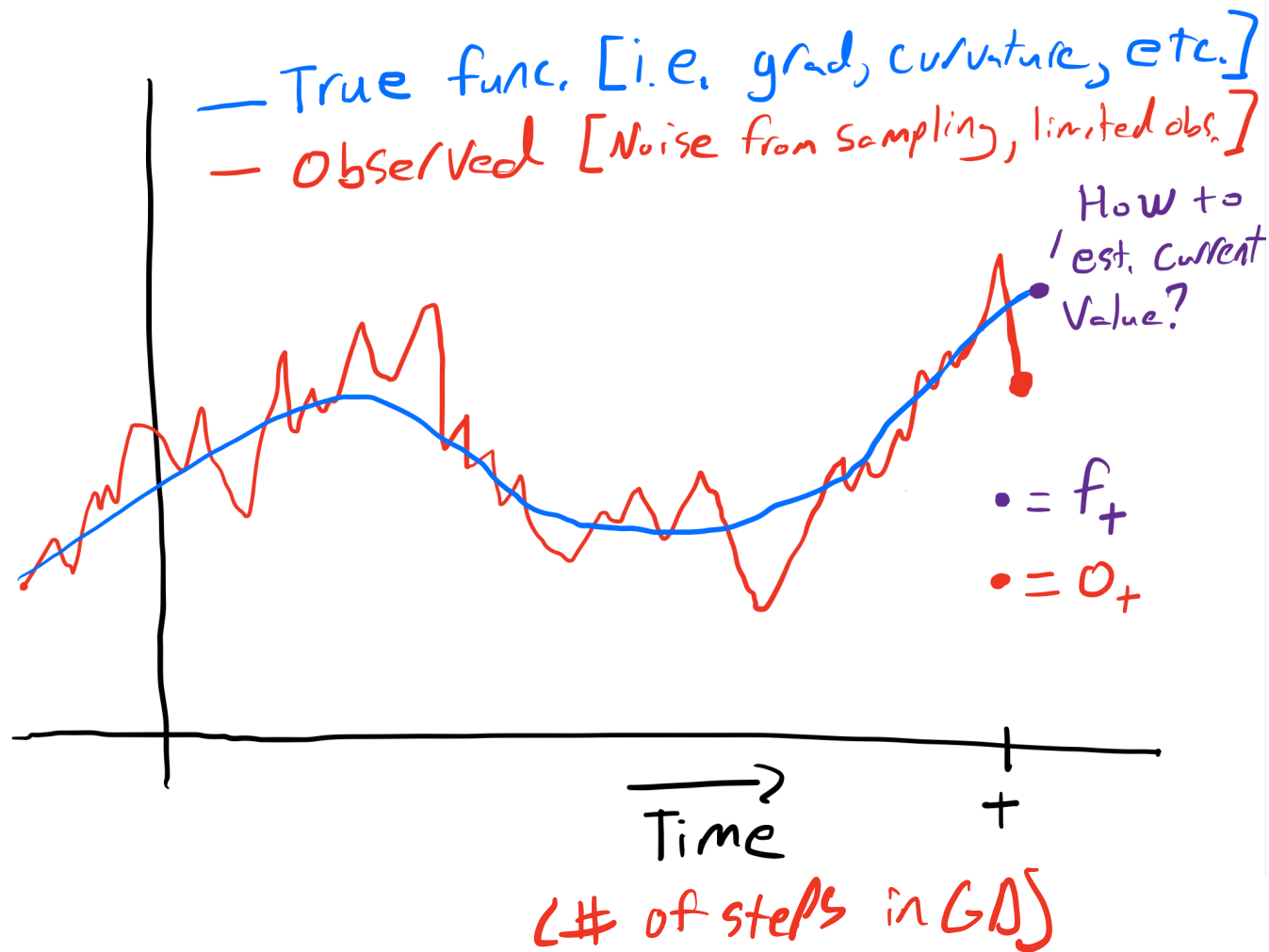
Stuck at saddle points or shallow optima

- Use momentum

Inconsistent scaling of the gradient

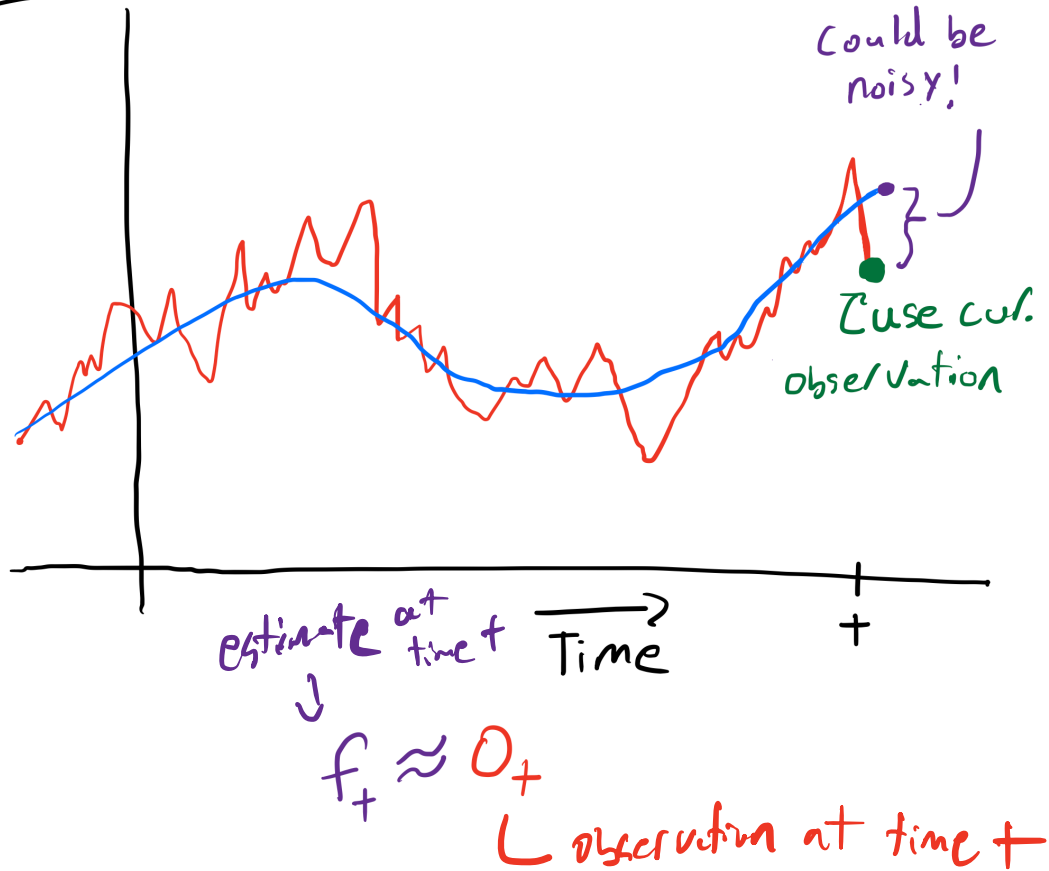
- Use RMSProp scaling

Exponential moving average (EMA)



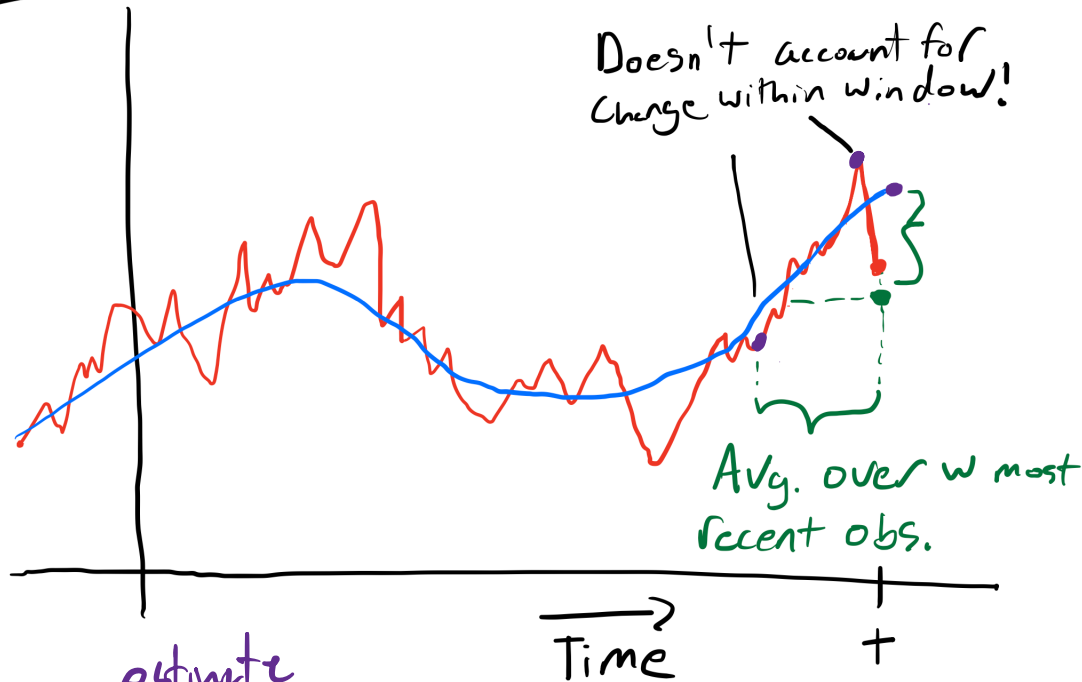
Exponential moving average (EMA)

Idea 1



Exponential moving average (EMA)

Idea 2 [Box avg.]



estimate

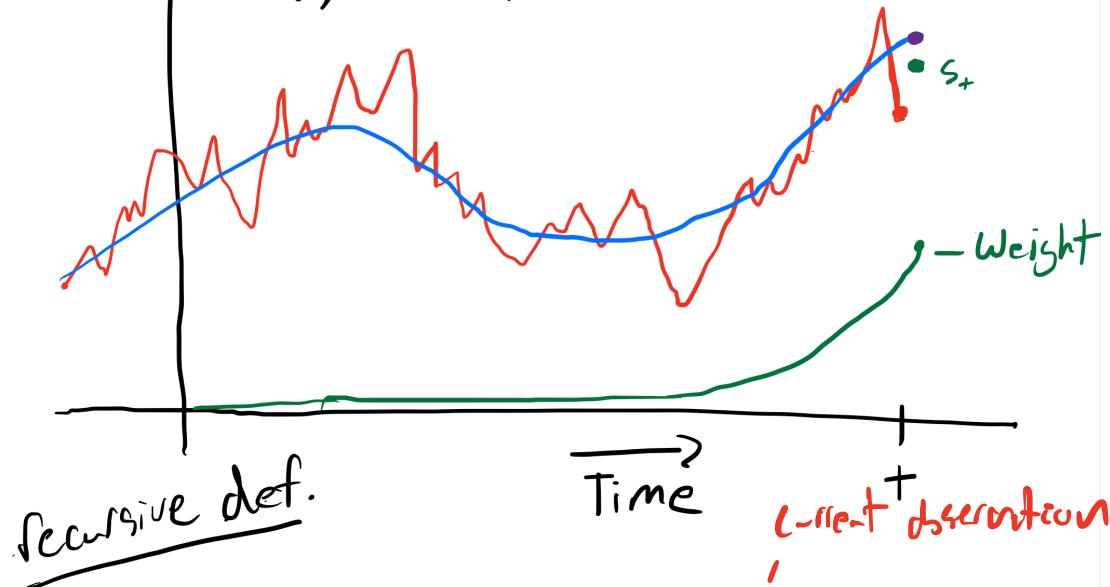
$$\hat{f}_t \approx \frac{1}{w} \sum_{i=t-w}^t o_i$$

Average obs. over a window of size w

Exponential moving average (EMA)

Idea 3 [EMA]

Average over entire history weighted by recency!



$$f_t \approx S_t$$

estimate at time t

$$S_t = (1-\alpha)O_t + \alpha S_{t-1}$$

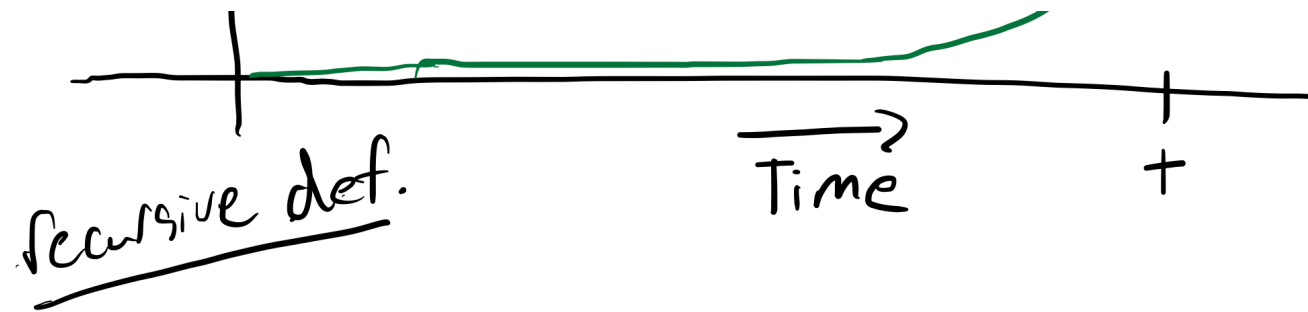
prev. at t-1

$$\sum_{t=0}^T \text{weight} = 1$$

$$\sum_{t=0}^{T-1} \text{weight} = \alpha$$

$$\alpha \in [0, 1]$$

Exponential moving average (EMA)



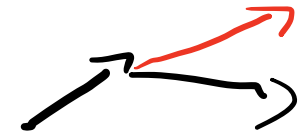
$$f_t \approx S_t \quad S_t = (1-\alpha)O_t + \alpha S_{t-1}$$

$$\rightarrow f_t \approx (1-\alpha)O_t + (1-\alpha)\alpha O_{t-1} + (1-\alpha)\alpha^2 O_{t-2} + \dots$$

Ex. $\alpha = 0.9$

$$f_t \approx 0.1 O_t + 0.09 O_{t-1} + 0.081 O_{t-2} + 0.0729 O_{t-3} + \dots$$

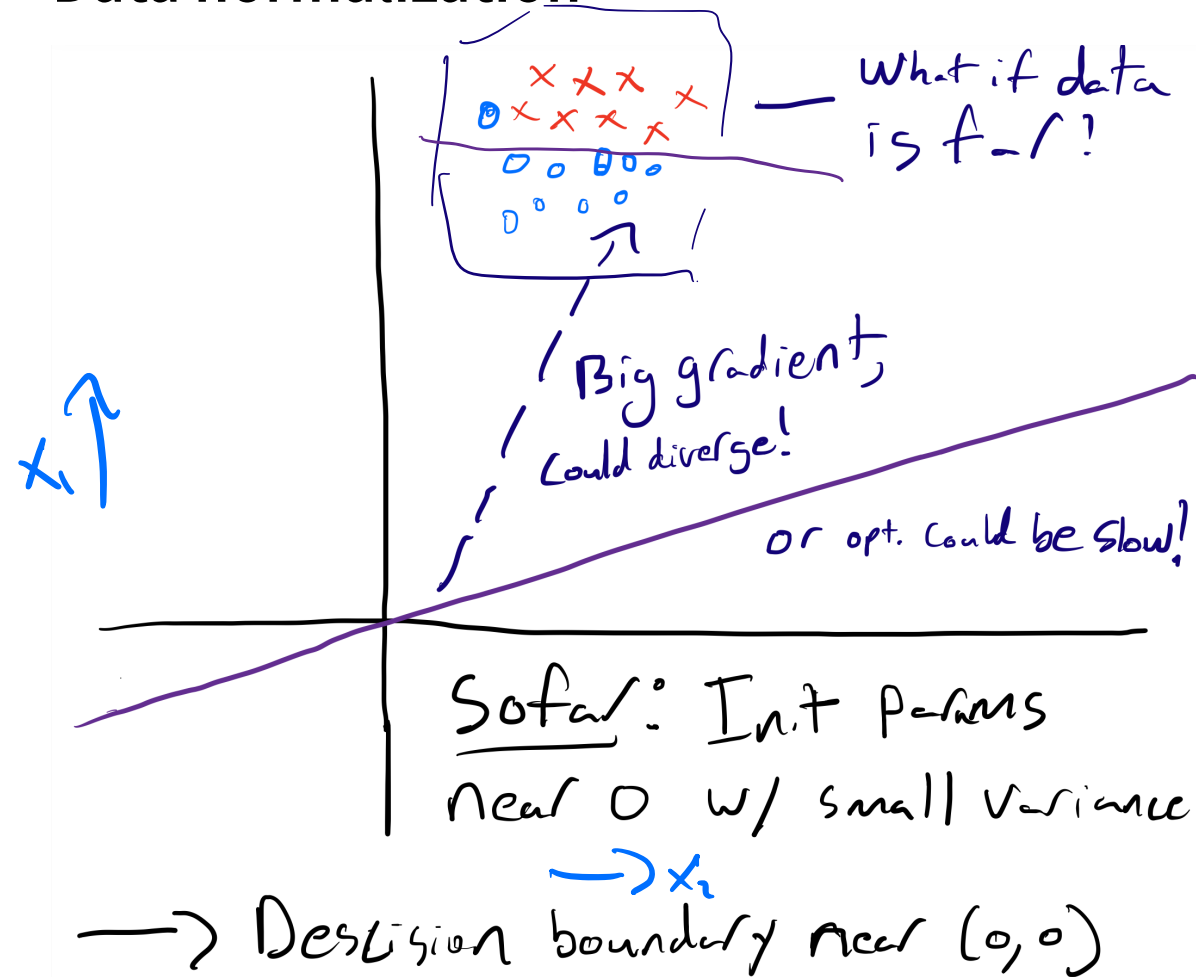
momentum



$$S_{t+1} = (1-\alpha)O_{t+1} + \alpha S_t$$

$$+ \boxed{\approx 0} O_{t-100}$$

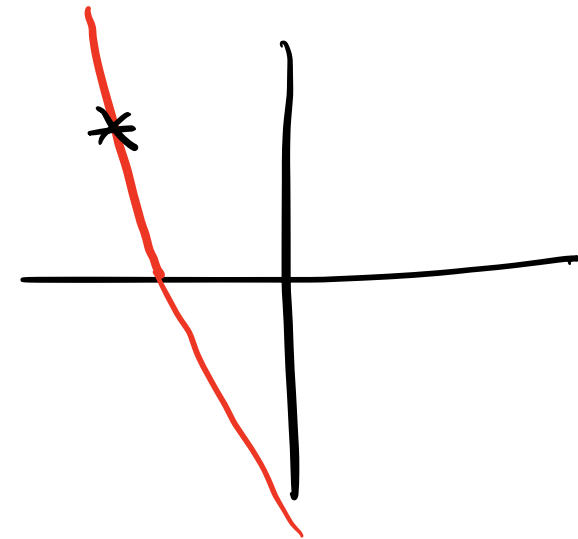
Data normalization



$$f(x) = \sum_i \underbrace{x_i w_i}_{\text{What about}}$$

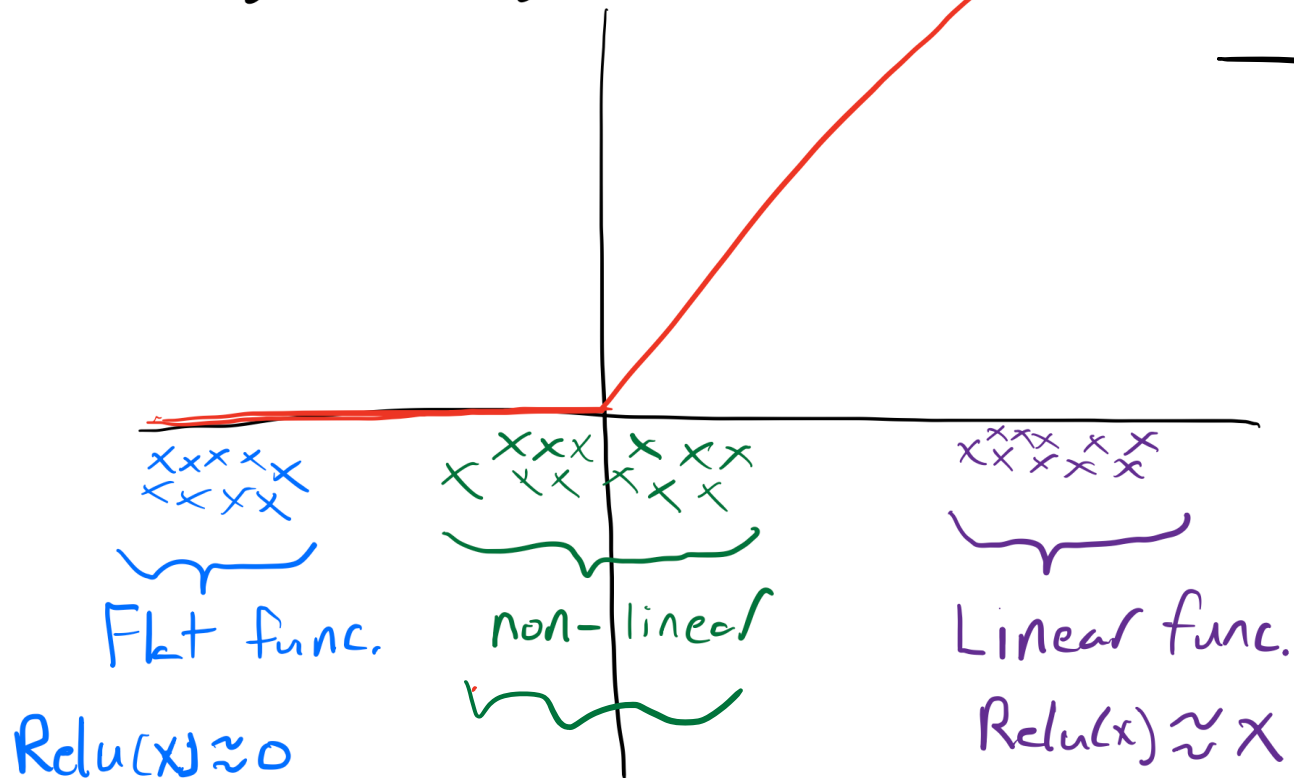
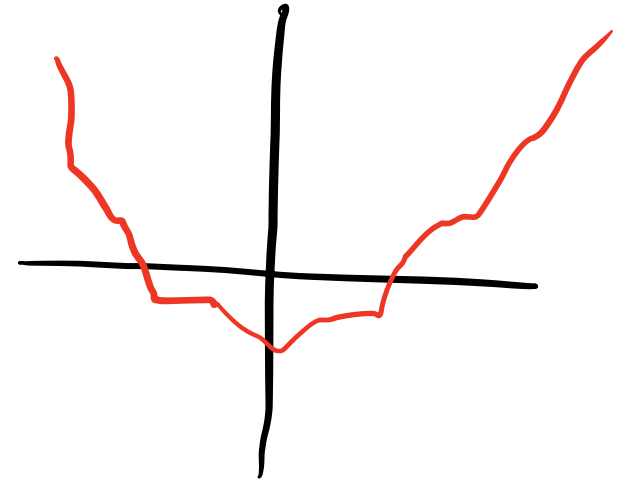
What about

~ decision boundary
 $f(x) = 0$

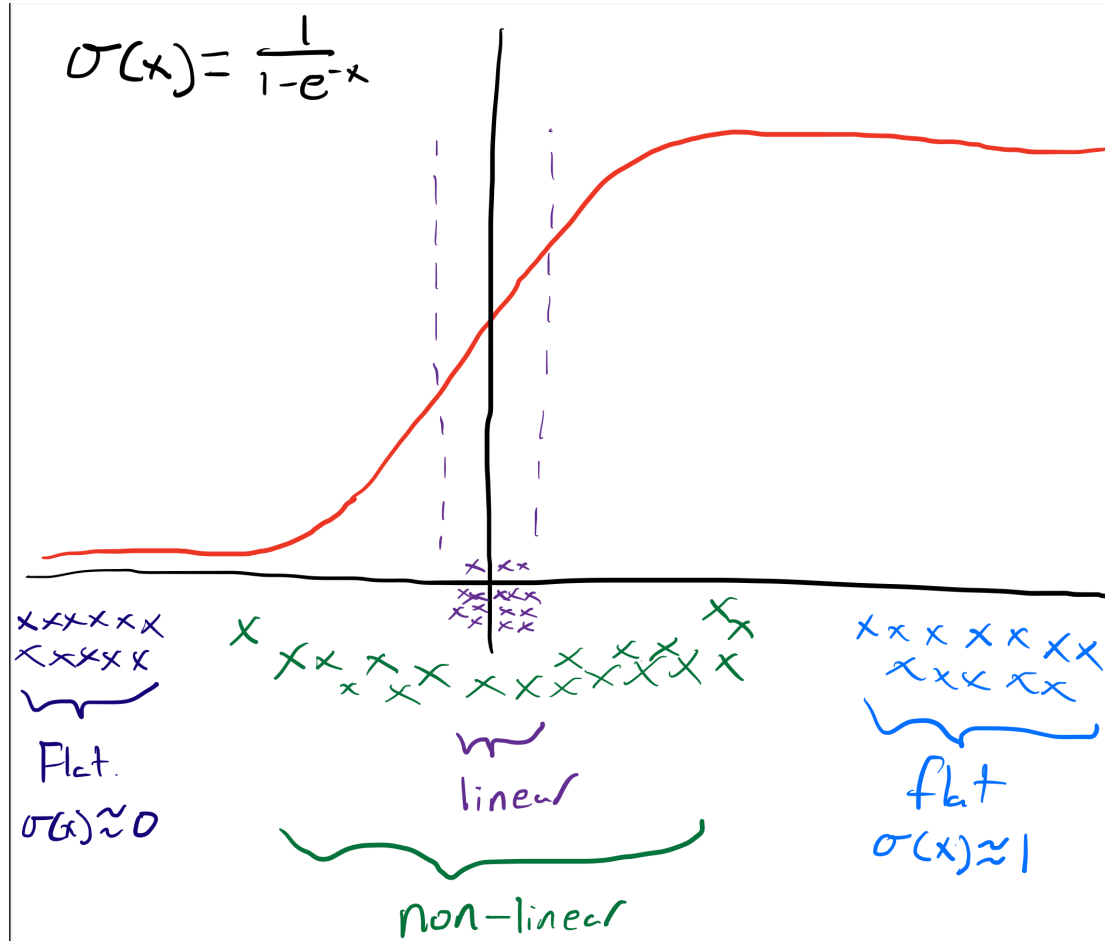


Data normalization

$$\text{Relu}(x) = \max(0, x)$$



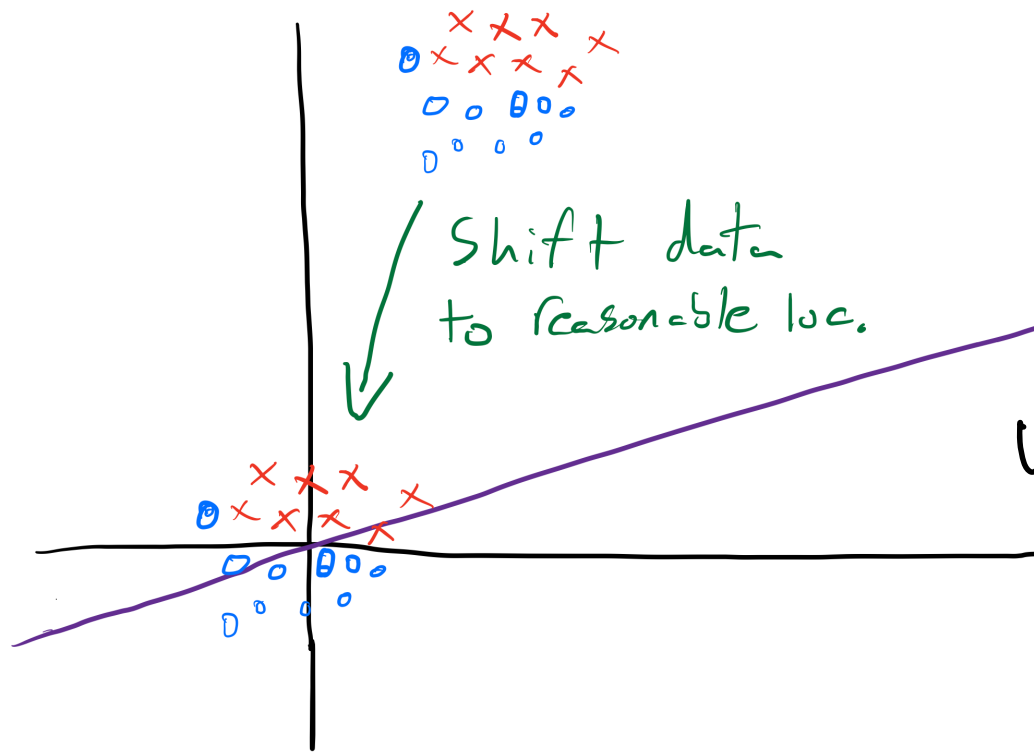
Data normalization



Data normalization



Data normalization



Decision boundary near (0,0)

$$X_1, X_2, \dots, X_N$$

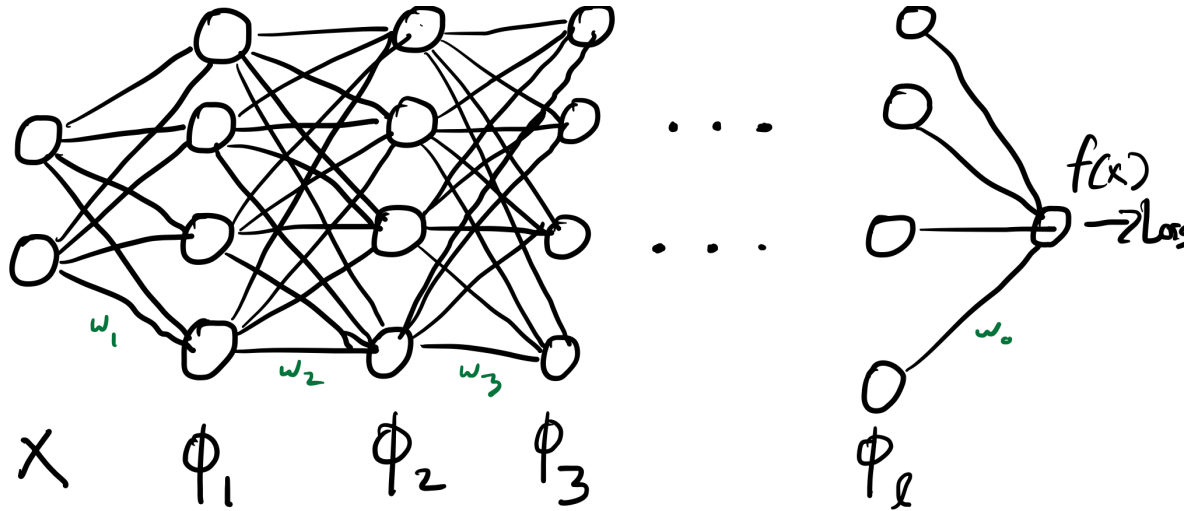
$$\frac{1}{N} \sum_{i=1}^N X_i \approx 0 \quad \leftarrow \text{want}$$

$$\sqrt{\frac{1}{N} \sum_{i=1}^N (x_i - \bar{x})^2} \approx 1 \quad \leftarrow \text{want}$$

Normalization

$$X_i \rightarrow \frac{X_i - \frac{1}{N} \sum_{j=1}^N X_j}{\sqrt{\frac{1}{N} \sum_{i=1}^N (x_i - \bar{x})^2}}$$

Vanishing and exploding gradients



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

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

$$\phi_3 = \sigma(\phi_2^T w_3 + b)$$

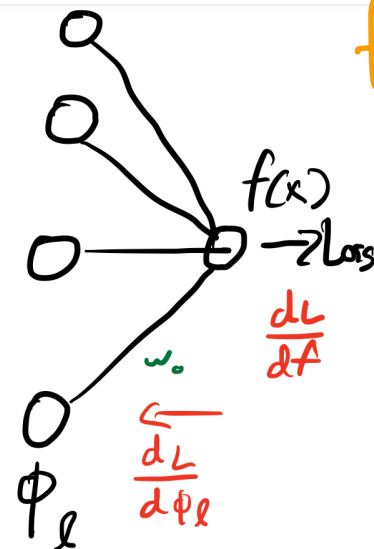
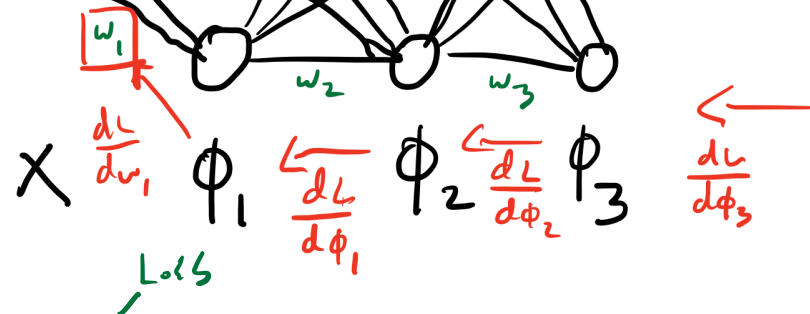
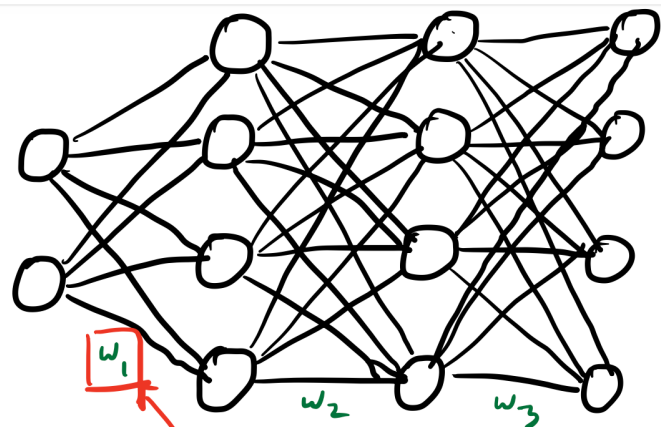
$\vdots$

$$f = \phi_l^T w_l + b \quad L = \text{Loss}(f)$$

l layers

many layers

Vanishing and exploding gradients



$$f_2(f_1(x)) = \frac{df_2}{df_1} \cdot \frac{df_1}{dx}$$

$$2, -5, 3, -6, \dots$$

$$= 180$$

$$\frac{dL}{dw_1} = \frac{d\phi_1}{dw_1} \cdot \frac{d\phi_2}{d\phi_1} \cdot \frac{d\phi_3}{d\phi_2} \cdots \frac{d\phi_l}{d\phi_{l-1}} \cdot \frac{df}{d\phi_l} \cdot \frac{dL}{df}$$

$$\frac{d}{dw_1} \sigma(xw_1 + b) = x \sigma'(xw_1 + b)$$

$$\text{for all } \frac{d\phi_i}{d\phi_{i-1}}$$

$$\frac{dL}{dw_1} = x \sigma'(xw_1 + b) \left(\prod_{i=2}^l \frac{d\phi_i}{d\phi_{i-1}} \right) \frac{df}{d\phi_l} \cdot \frac{dL}{df}$$

$$\left| \frac{d\phi_i}{d\phi_{i-1}} \right| \approx M$$

Scale

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

Vanishing and exploding gradients

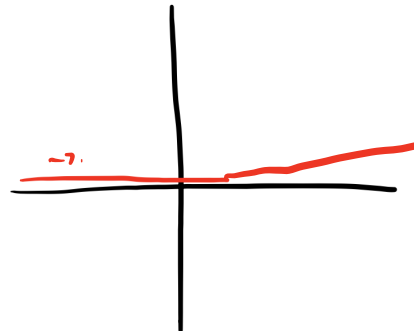
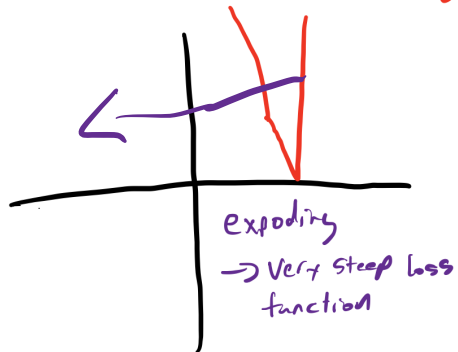
$$\text{If } M > 1 \rightarrow \left| \frac{dl}{dw_i} \right| \gg 1$$

Gradient explodes!

$$\text{If } M < 1 \rightarrow \left| \frac{dl}{dw_i} \right| \approx 0$$

Gradient Vanishes

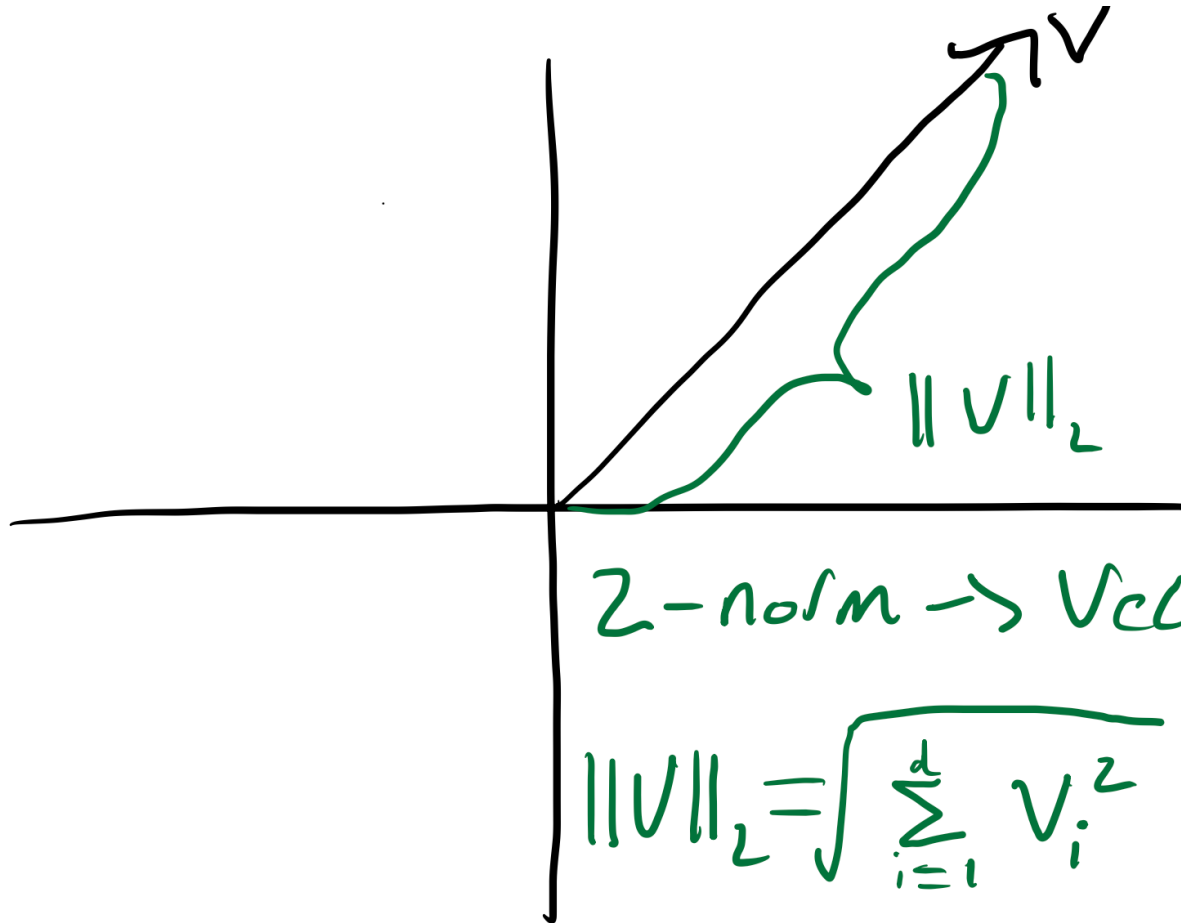
Both bad for gradient descent!



$$\approx M^x$$

$$9^{20} \approx 0$$
$$2^{20}$$

Gradient clipping



Gradient clipping

Explicitly clip the gradient to prevent it from becoming too large.

$\epsilon = \text{limit}$

clip by value

$$\text{clip}_{\text{value}}(\mathbf{x}, \epsilon) = \begin{bmatrix} \min(\max(x_1, -\epsilon), \epsilon) \\ \min(\max(x_2, -\epsilon), \epsilon) \\ \vdots \end{bmatrix}$$

eg. gradient

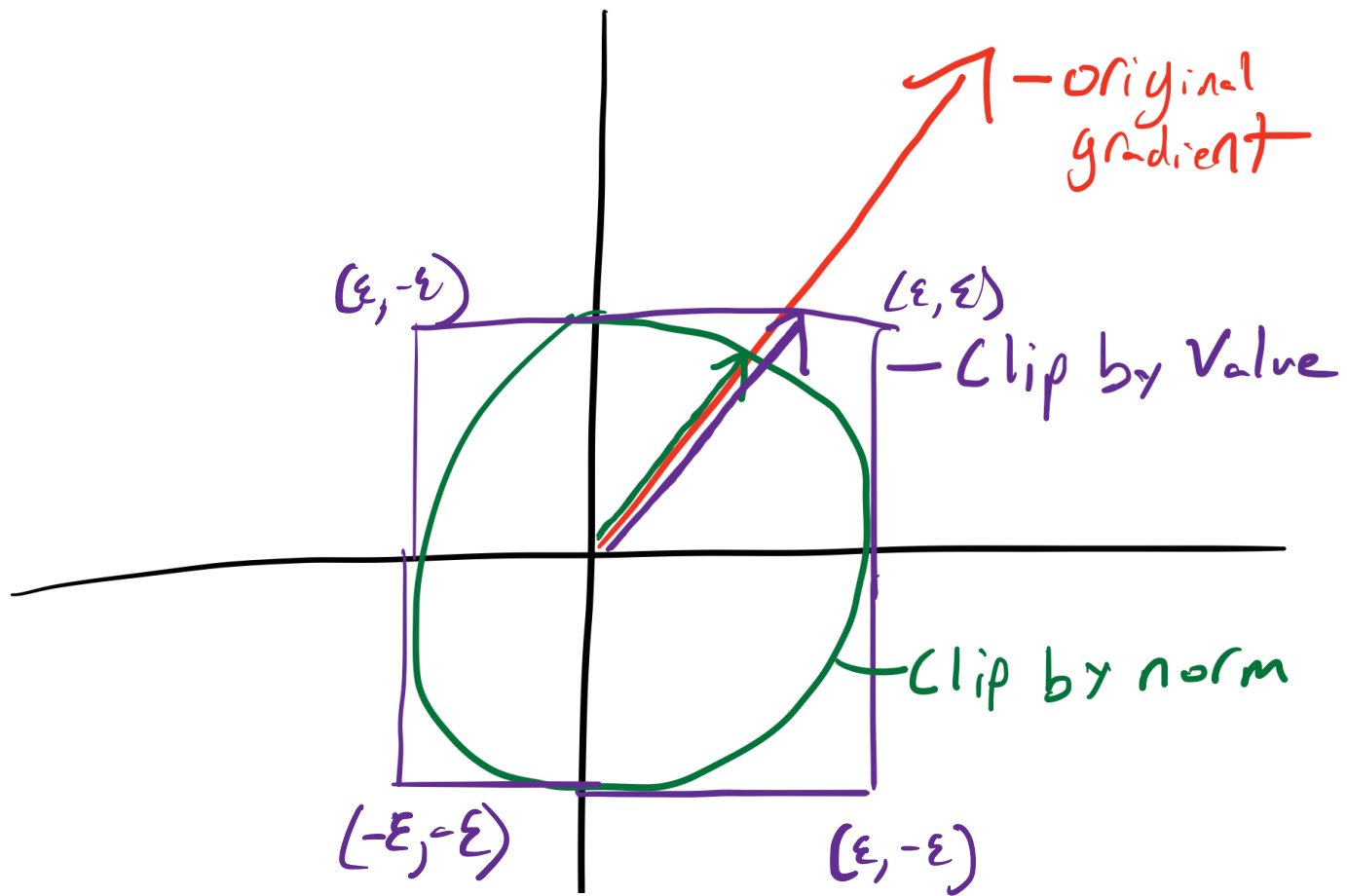
enforce that $\mathbf{x}$ has length ϵ

$$\text{clip}_{\text{norm}}(\mathbf{x}, \epsilon) = \begin{cases} \frac{\epsilon \mathbf{x}}{\|\mathbf{x}\|_2} & \text{if: } \|\mathbf{x}\|_2 > \epsilon \\ \mathbf{x} & \text{if: } \|\mathbf{x}\|_2 \leq \epsilon \end{cases}$$

Clipped
gradient
descent

$$\mathbf{w}^{(k+1)} \leftarrow \mathbf{w}^{(k)} - \alpha \text{clip}(\nabla_{\mathbf{w}} \text{Loss}(\mathbf{w}^{(k)}, \mathbf{X}, \mathbf{y}))$$

Gradient clipping



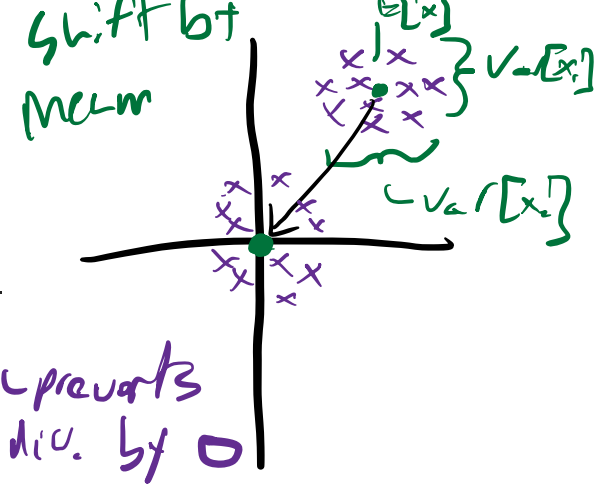
Normalization

Batch normalization

Normalize over the batch:

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

observation

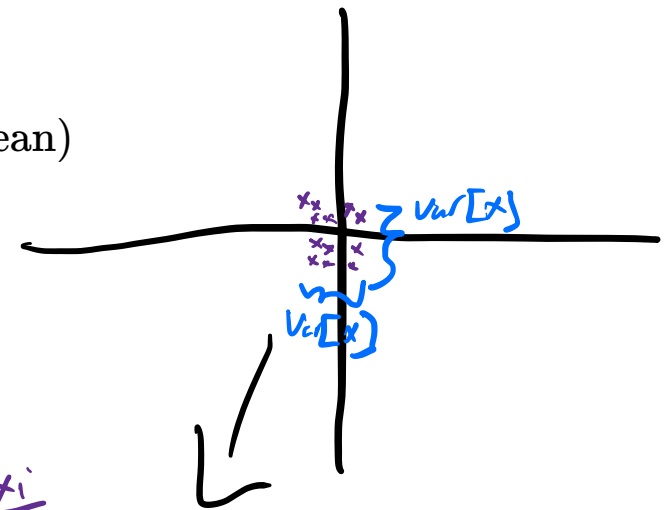


Training time:

Batch: $\{x_1, x_2, \dots, x_B\}$

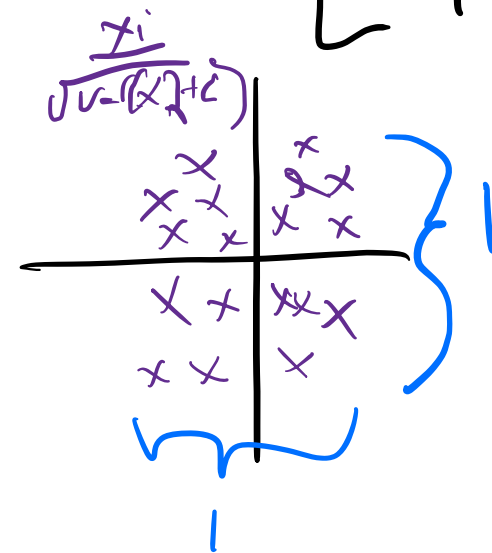
$$\mathbb{E}[x] \approx \bar{x} = \frac{1}{B} \sum_{i=1}^B x_i \quad (\text{sample mean})$$

out current estimate of the mean $\mathbb{E}[x]$



note

Could also use full dataset



Batch normalization

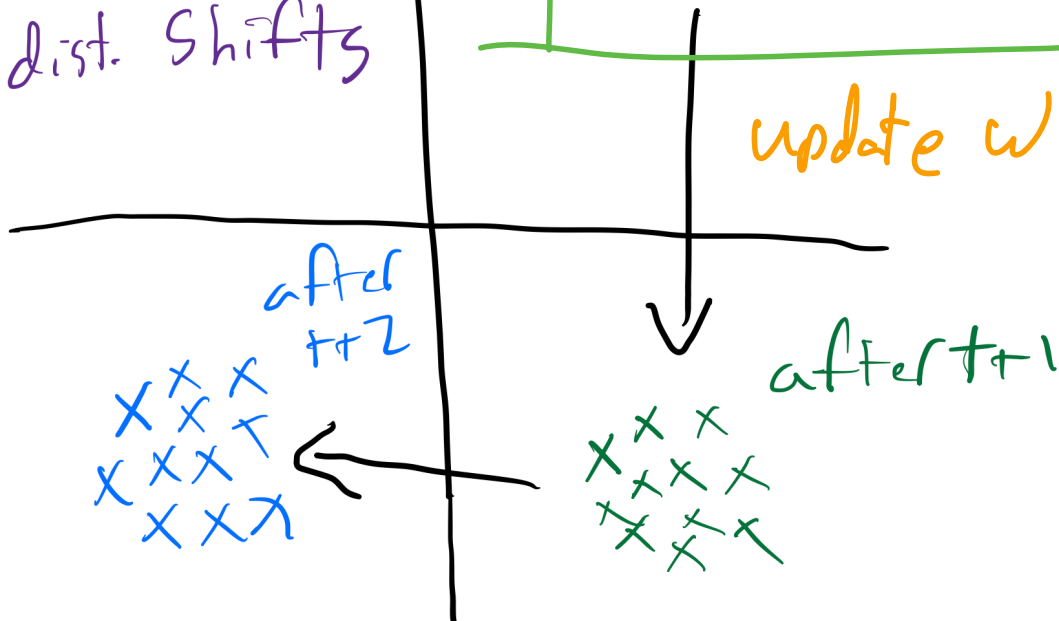
- Layer i

$$\phi_i(xw)$$

As weights change
data dist. shifts



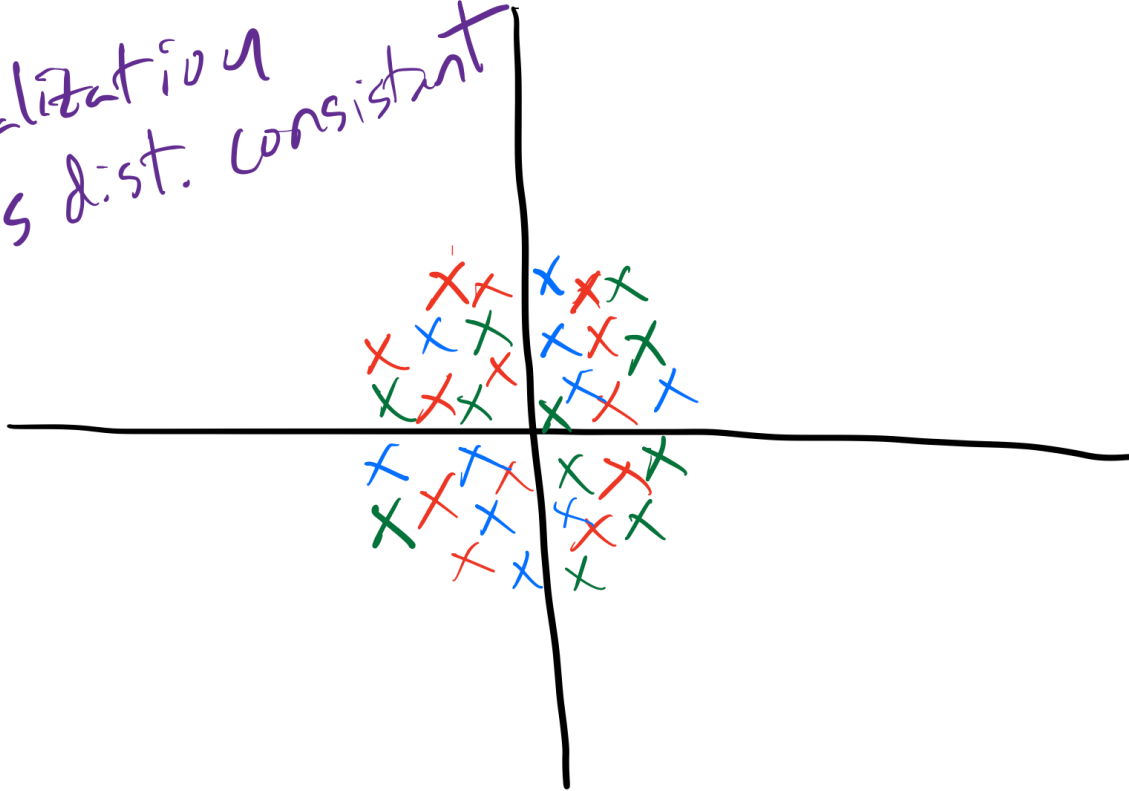
update w



Batch normalization

w/ normalization

normalization
keeps dist. consistent



Batch normalization

Biased estimator:

$$\text{Var}[x] \approx s^2 = \frac{1}{B} \sum_{i=1}^B \left(x_i - \left(\frac{1}{B} \sum_{i=1}^B x_i \right) \right)^2 \quad (\text{sample var.})$$

Sample mean

Unbiased estimator:

$$\text{Var}[x] \approx s^2 = \frac{1}{B-1} \sum_{i=1}^B \left(x_i - \left(\frac{1}{B} \sum_{i=1}^B x_i \right) \right)^2 \quad (\text{sample var.})$$

$$\text{BatchNorm}_{\text{train}}(x) = \frac{x - \bar{x}}{\sqrt{s^2 + \epsilon}}$$

Batch normalization

Running estimate:

$$\bar{\mu}^{(k+1)} \longleftarrow \beta \bar{\mu}^{(k)} + (1 - \beta) \bar{x}^{(k)}$$

$$\bar{\sigma}^{2(k+1)} \longleftarrow \beta \bar{\sigma}^{2(k)} + (1 - \beta) s^{2(k)}$$

$$\text{BatchNorm}_{\text{test}}(x) = \frac{x - \bar{\mu}}{\sqrt{\bar{\sigma}^2 + \epsilon}}$$

Layer normalization

Normalize over the layer:

$$\text{LayerNorm}(\mathbf{x}) = \frac{\mathbf{x} - \bar{x}}{\sqrt{s^2 + \epsilon}}, \quad \mathbf{x} = \begin{bmatrix} x_1 \\ \vdots \\ x_d \end{bmatrix}$$

Training & test time:

$$\bar{x} = \frac{1}{d} \sum_{i=1}^d x_i \quad (\text{output mean})$$

Layer normalization

Biased estimator:

$$s^2 = \frac{1}{d} \sum_{i=1}^d \left(x_i - \left(\frac{1}{d} \sum_{i=1}^d x_i \right) \right)^2 \quad (\text{output var.})$$

Unbiased estimator:

$$s^2 = \frac{1}{d-1} \sum_{i=1}^d \left(x_i - \left(\frac{1}{d} \sum_{i=1}^d x_i \right) \right)^2 \quad (\text{output var.})$$

Scaled normalization

$$\text{BatchNorm}(x) = \frac{x - \mathbb{E}[x]}{\sqrt{\text{Var}[x] + \epsilon}} \gamma + \kappa$$

$$\text{LayerNorm}(\mathbf{x}) = \frac{\mathbf{x} - \bar{x}}{\sqrt{s^2 + \epsilon}} \gamma + \kappa$$