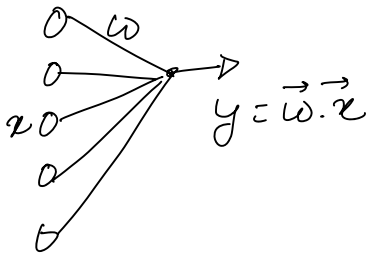


Pruning a neuron.

$$y = (w, x)$$

$$w, x \in \mathbb{R}^d$$



Q. What should we try to preserve when pruning?

- Behavior of neuron.
- What function it was representing

Remark: While pruning at "Inference" time the considerations are different - we want to preserve the function itself as if we were thinking about training a model and its expressive power matters.

To model behavior: we want to say preserve the following under compression

If $\hat{w} \in \mathbb{R}^d$ s.t. $\|\hat{w}\|_1 \leq K$
we want to preserve the f^n

$$\hat{w} \cdot x \simeq w \cdot x \quad \forall x$$

Let X be the n representative data
then

$$\hat{w} = \underset{\|\hat{w}\|_1 \leq \kappa}{\operatorname{argmin}} \left[\|\hat{w} \cdot x - w \cdot x\|_2 \right]$$

Write this as,

$$\hat{w}, M = \underset{\substack{\hat{w} \in \mathbb{R}^d \\ M \in [0,1]^d \\ \|M\|_1 \leq \kappa}}{\operatorname{argmin}} \left[\left\| (\hat{w} \odot M) \cdot x - w \cdot x \right\|_2^2 \right]$$

Jointly optimising M & $\hat{w}$
is NP Hard (Blum et al & Davies
2008)

$\rightarrow$ Greedy Algorithm

Which element to prune first?

→ Magnitude? $|w_i| \uparrow$ but $|x_i|$

→ $f(|w_i|, |x_i|)$? very low

→ does not
lose the
Correlations.

$$y = w \cdot x$$

↑
true model

$$\hat{y} = \hat{w} \cdot x$$

↑
approximate model.

$$\min \| \hat{w} \cdot x - w \cdot x \|_2^2$$

$$\| \hat{w} \| \leq d-1$$

$$\hat{w}_i = 0 \text{ for some } i$$

$$\min_{w_1=0} \| \hat{w} \cdot x - w \cdot x \|_2^2$$

$$\min_{w_2=0} \| \hat{w} \cdot x - w \cdot x \|_2^2$$

⋮

$$\min_{w_d=0} \| \hat{w} \cdot x - w \cdot x \|_2^2$$

→ whichever
gives the
best loss
we choose
that "i"

Subproblem.

$$(i) \quad \hat{\omega}^{(i)} = \underset{\omega}{\operatorname{argmin}} \left\| \underbrace{X_{-i}^T \omega}_{\text{linear pred.}} - \underbrace{y_i}_{\gamma} \right\|^2$$

$$\left[\begin{array}{c|c} X & \text{shaded column} \end{array} \right] \quad \left| \begin{array}{c} y_i \\ \leftarrow i \\ = y \end{array} \right.$$

Normal eqs. $\rightarrow X_{-i}^T \hat{\omega} = y_i$

$$\hat{\omega} = (X_{-i}^T X_{-i})^{-1} \cdot X_{-i}^T y_i$$

loss: $\hat{y} = X_{-i}^T \hat{\omega}$

$$\underline{\underline{L_i}} = \left(\left\| \hat{y} - y_i \right\|_2^2 \right) \leftarrow \begin{array}{l} \text{choose} \\ \text{that} \\ \text{least} \end{array}$$

① "Saliency"

② update $\omega_{-i} \leftarrow \hat{\omega}$

Pruning a Moel