

Let F be the model.; $\vec{w}$ be the flattened parameter vector, X : input data representative of input

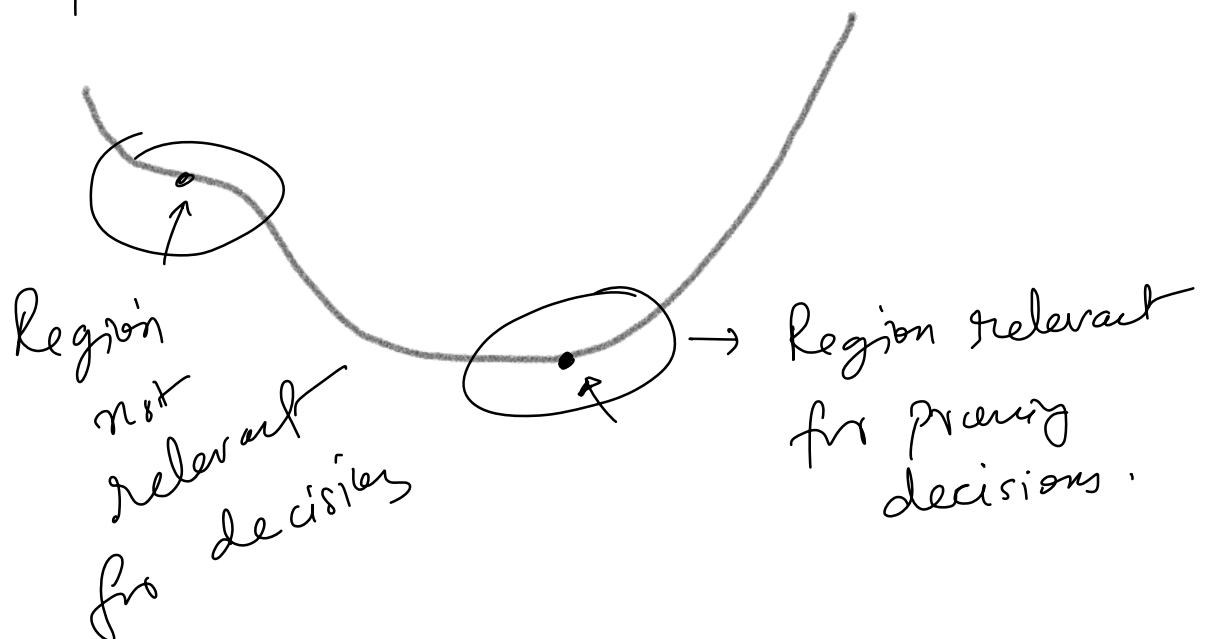
Q Again, what should we be looking at while pruning?

The only "information" that we get about good and bad decisions in ML are through the loss landscape,

↳ "Training depends on the loss landscape"

Q: Which Model should we start at?

→ Since the information you can get about loss landscape is "local", you better start from the "Converged Model":



Let $L(w^*, X)$ be the loss function and let w^* be the weight at minimum (converged model) then to get "information" about the landscape analytically we need to approximate the local landscape

$$L(w, X) = L(w^*, X) + \nabla L(w^*, X)^T \delta w + \frac{1}{2} \delta w^T \nabla^2 L(w^*, X) \delta w + O(\|\delta w\|_2^3)$$

→ "converged" $\Rightarrow \nabla L(w^*, X)^T = 0$

"local" / "quadratic" $\Rightarrow O(\|\delta w\|_2^2)$

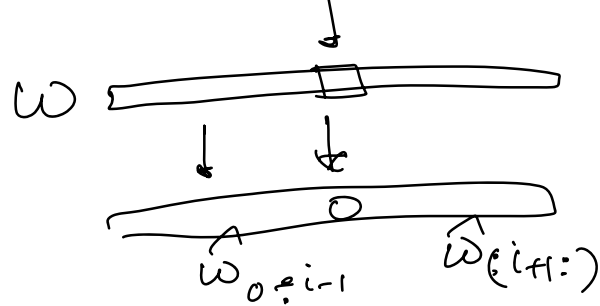
$$L(w, X) - L(w^*, X) = \frac{1}{2} \delta w^T H(w^*, X) \delta w$$

$$E(\delta w) = \frac{1}{2} \delta w^T H(w^*, X) \delta w$$

Note that this is just a quadratic form and thus is easy to analyse

which "i" to prune

- ↳ lowest saliency
- ↳ update the rest of the wts.



N: Subproblems

if "i":

$$\delta w = \underset{\delta w_i = -w_i}{\text{argmin}} [E(\delta w)]$$

So that the i^{th} wt goes to 0

Lagrange's multiplier method.

$$\delta w = \underset{\delta w}{\text{argmin}} \left[\underbrace{E(\delta w)}_{\frac{1}{2} \delta w^T H \delta w} + \lambda \left(\underbrace{e_i^T (w + \delta w)}_{\text{unit vector}} \right) \right]$$

$$\frac{\partial T}{\partial \lambda} = 0 \Rightarrow \boxed{\delta w_i = -w_i}^T$$

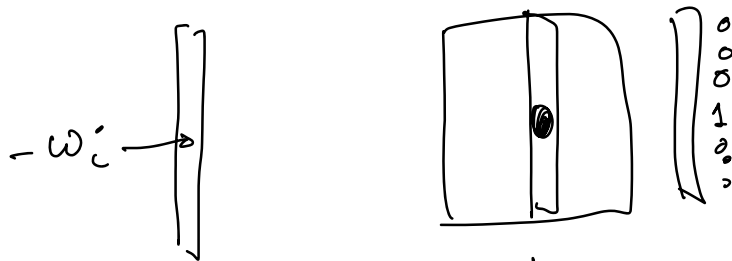
|| for matrix operations look at matrix cookbook.

$$\frac{\partial T}{\partial (\delta w)} = H \delta w + \lambda e_i = 0$$

often loc
in $\delta\omega$
is known

$$H\delta\omega + \lambda e_i = 0$$

$$\delta\omega \neq \lambda H^{-1} e_i = 0$$



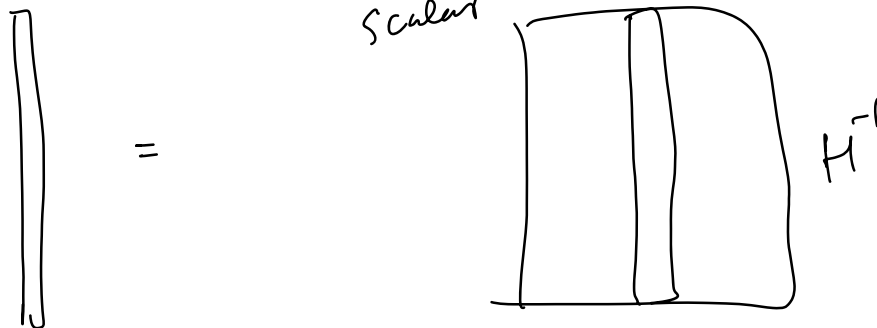
$$-w_i = \lambda H^{-1} e_i$$

$$\therefore \lambda = \frac{-w_i}{[H^{-1}]_{ii}}$$

$$\delta\omega = -\lambda H^{-1} e_i$$

$$\delta\omega = \left(\frac{w_i}{[H^{-1}]_{ii}} \right) H^{-1} e_i$$

↑ scalar
↑ ith column



Saliency :

$$\frac{1}{2} \delta\omega^T H \delta\omega$$

$$= \frac{1}{2} \frac{w_i}{(H^{-1})_{ii}^2} \left[(H^T e_i)^T \underbrace{H H^T}_{I} e_i \right]$$

$$= \frac{1}{2} \left[\frac{w_i^2}{(H^{-1})_{ii}^2} \right] (e_i^T \underbrace{(H^{-1})^T}_{I} e_i)$$

$$= \frac{1}{2} \frac{w_i^2}{H_{ii}^{-1}{}^2} \cdot H_{ii}^{-1}$$

$$\boxed{\text{Saliency} = \frac{1}{2} \frac{w_i^2}{H_{ii}^{-1}{}^2}}$$

$$\boxed{\delta w = \frac{w_i}{H_{ii}^{-1}} H^T e_i}$$

* Algorithm :

① Train the full Model.

② $(S_i \leftarrow \text{Saliency of } i)$

pick i with least saliency.

update

$w \leftarrow w + \delta w$.

if $\{S_i\}_{\min}$ is large compared to loss

then Retrain

Redo
local
approx

$g(\{s_i\}_{i=1}^n)$ is small

Complexity - :

1 prime :

$O(N^3)$