

Structured Pruning of a Model under the constraint that we can only have enough memory for inference.

Q: How do we prune?

A: SparseGPT, Wanda, Magnitude $\rightarrow$ Yes we can use the saliency and do a one shot pruning. But can we do better?

* We can use reasonable amount of forward passes.

* Idea 1: Sample some $\hat{M}_i$ (pruned Models) and test it on downstream tasks.

Q1: How to Sample

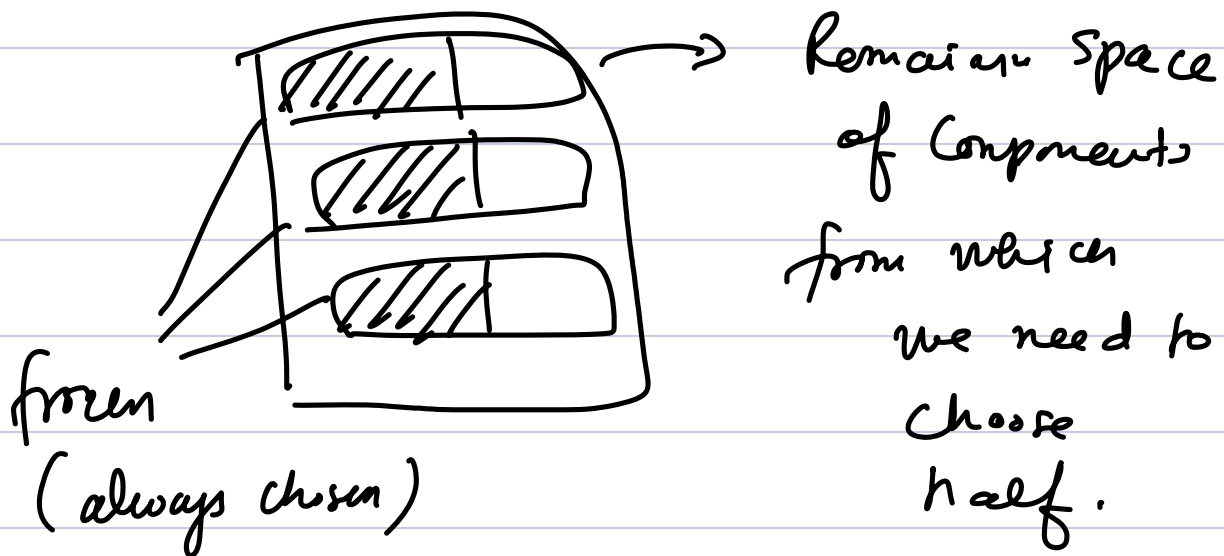
Q2: What do I do with the data $\{\hat{M}_i, S_i\}$ $\hat{M}_i$:

Sampling.

* Space is too large, there are some obvious Wrong Samples. For ex. if you don't sample any thing in a layer that's a broken model so let us keep some models in each layer

* How do you choose the Components to keep?
Lds use "Saliency" from existing methods

In fact if my target sparsity is $p \in (0, 1)$
let us keep $(1 - 2p)$ fraction of Components fixed.



* How to Sample? Uniform?

→ Use saliency to sample with some prior.

② What do I do with $\{\hat{M}_i, s_i\}$

→ Just pick best (one naive way)

Q Can we do better?

possible modules (neurons, heads, etc) are much larger than "reasonable # of forward passes" which means that we won't be able to get information on most configurations / modules if done naively.

Q. What do we do?

Idea: Create a Model of utility

$$s_i = \hat{U}(\hat{M}_i)$$

↓

approx utility

$$\hat{s}_i = \hat{U}(\{m_1, m_2, \dots, m_n\})$$

What is the Simplest Model?

linear model

$$\hat{S}_i = \sum_{m_j \in M_i} \hat{\beta}_j(m_j)$$

N : # models

linear Regression problem

n : # data
Samples
or
forward
passes

$$\hat{M} \begin{bmatrix} 0 & + & 1 & - & 0 & + \\ 1 & 1 & 1 & 0 & 0 & 1 \end{bmatrix} \begin{bmatrix} \beta \end{bmatrix} = \begin{bmatrix} \vec{S}_i \end{bmatrix}$$

$n \times N$ $N \times 1$ $n \times 1$ $N \gg n$

Sparse Matrix X

We want to estimate β .

Once we get β and if the model Really holds then all we need to do is choose $\sim$ half (i.e. p fraction of entire model) models with highest " β " scores.

$$\hat{M} \beta = S \leftarrow \text{underspecified}$$

$$\beta^* = \arg \min || \hat{M} \beta - s ||_2 + \underbrace{\gamma || \beta ||}_2$$

Thus add regularisation

Q: How do you sample matrix ($\hat{M}$)

Q $\rightarrow$ prior based yes. but how many non-zeros per row?

A: $\hookrightarrow$ try to maintain fraction "p"

Q: Even with prior it is likely that we miss some module altogether.

$\Rightarrow \beta_i$ for that module would be a BAD estimate

Trick: If you sample $\hat{M}_i$ also sample its complement, i.e. $\hat{M}_i^c$

$$\hat{M}_i^c = \{ m_i \mid m_i \notin \hat{M}_i \text{ and}$$

it comes from set of modules under consideration. }

Exercise: Prove that having the complement of a sample in the data helps reduce variance of β estimation.