

Taylor Expansion of LogSumExp and Softmax

Su Jianlin

2026-07-13

Recently, I saw the paper *The Key to Going Linear: Analysis-Driven Transformer Linearization*, which directly performs an approximate expansion on Softmax to linearize Attention. I gained some insights from it and will record them here.

Some Notations

First, we introduce the following notations:

$$\mathbf{x} = (x_1, x_2, \dots, x_n) \in \mathbb{R}^n, \quad \bar{\mathbf{x}} = \frac{1}{n} \sum_{i=1}^n x_i \quad (1)$$

$$\text{logsumexp}(\mathbf{x}) = \log \sum_{i=1}^n e^{x_i} = \log n + \log \bar{e^{\mathbf{x}}} \quad (2)$$

$$\text{softmax}(\mathbf{x}) = \frac{e^{\mathbf{x}}}{\sum_{i=1}^n e^{x_i}} = e^{\mathbf{x} - \text{logsumexp}(\mathbf{x})} \quad (3)$$

We agree that function operations on vectors are in the Hadamard sense, for example

$$e^{\mathbf{x}} = (e^{x_1}, e^{x_2}, \dots, e^{x_n}), \quad \mathbf{x}^2 = (x_1^2, x_2^2, \dots, x_n^2) \quad (4)$$

In particular, carefully note the difference between $\bar{e^{\mathbf{x}}}$ and $e^{\bar{\mathbf{x}}}$, as well as $\overline{\mathbf{x}^2}$ and $\bar{\mathbf{x}}^2$: the former applies the function to each component first and then averages, while the latter averages first and then applies the function. The two are generally not equal.

In addition to the above definitions, logsumexp and softmax can also be connected through their gradients:

$$\text{softmax}(\mathbf{x}) = \nabla_{\mathbf{x}} \text{logsumexp}(\mathbf{x}) \quad (5)$$

Basic Expansion

Next, we expand $\text{logsumexp}(\mathbf{x})$. Introducing $\mathbf{y} = \mathbf{x} - \bar{\mathbf{x}}$, we have

$$\begin{aligned}
\text{logsumexp}(\mathbf{x}) &= \log n + \bar{\mathbf{x}} + \log \bar{e^{\mathbf{y}}} \\
&= \log n + \bar{\mathbf{x}} + \log \left(1 + \mathbf{y} + \frac{\mathbf{y}^2}{2} + \frac{\mathbf{y}^3}{6} + \frac{\mathbf{y}^4}{24} + \dots \right) \\
&= \log n + \bar{\mathbf{x}} + \log \left(1 + \frac{\bar{\mathbf{y}}^2}{2} + \frac{\bar{\mathbf{y}}^3}{6} + \frac{\bar{\mathbf{y}}^4}{24} + \dots \right) \\
&= \log n + \bar{\mathbf{x}} + \frac{\bar{\mathbf{y}}^2}{2} + \frac{\bar{\mathbf{y}}^3}{6} + \left(\frac{\bar{\mathbf{y}}^4}{24} - \frac{(\bar{\mathbf{y}}^2)^2}{8} \right) + \dots
\end{aligned} \tag{6}$$

The last step uses $\log(1+t) = t - t^2/2 + t^3/3 - \dots$. If needed, one can continue to expand further. If the reader is not familiar with this, they can also leave it to Kimi to complete.

Introducing the bias of $\bar{\mathbf{x}}$ here is a small trick for simplification. If it is not introduced, an equivalent result can still be obtained, but the various terms of $\mathbf{y}$ will be explicitly expanded:

$$\text{logsumexp}(\mathbf{x}) = \log n + \bar{\mathbf{x}} + \underbrace{\left(\frac{\bar{\mathbf{x}}^2}{2} - \frac{\bar{\mathbf{x}}^2}{2} \right)}_{\bar{\mathbf{y}}^2/2} + \underbrace{\left(\frac{\bar{\mathbf{x}}^3}{6} - \frac{\bar{\mathbf{x}}\bar{\mathbf{x}}^2}{2} + \frac{\bar{\mathbf{x}}^3}{3} \right)}_{\bar{\mathbf{y}}^3/6} + \dots \tag{7}$$

Finding the Gradient

For $\text{softmax}(\mathbf{x})$, we directly use the identity (5) and differentiate both sides of Equation (6) to obtain the expansion of $\text{softmax}(\mathbf{x})$. To do this, we utilize

$$\nabla_{\mathbf{x}} \bar{\mathbf{x}} = \frac{\mathbf{1}}{n}, \quad \nabla_{\mathbf{x}} \bar{\mathbf{y}}^m = \frac{m}{n} \left(\mathbf{y}^{m-1} - \bar{\mathbf{y}}^{m-1} \right) \tag{8}$$

to get

$$\text{softmax}(\mathbf{x}) = \frac{1}{n} \left(1 + \mathbf{y} + \left(\frac{\mathbf{y}^2}{2} - \frac{\bar{\mathbf{y}}^2}{2} \right) + \left(\frac{\mathbf{y}^3}{6} - \frac{\bar{\mathbf{y}}^3}{6} - \frac{\bar{\mathbf{y}}^2 \mathbf{y}}{2} \right) + \dots \right) \tag{9}$$

Note that if a finite number of terms are truncated, the right-hand side can still maintain the property that the sum of all components is 1 (which can be seen from the gradient form of $\nabla_{\mathbf{x}} \bar{\mathbf{y}}^m$), but it cannot guarantee the non-negativity of each component. This point needs to be noted in practical applications (for example, when taking the logarithm to calculate entropy).

Besides finding the gradient, we can also expand using $\text{softmax}(\mathbf{x}) = e^{\mathbf{x} - \text{logsumexp}(\mathbf{x})} = e^{\mathbf{y} - \text{logsumexp}(\mathbf{y})}$. Combining this with Equation (6), we get:

$$\begin{aligned} e^{\mathbf{y}} e^{-\text{logsumexp}(\mathbf{y})} &= \left(1 + \mathbf{y} + \frac{\mathbf{y}^2}{2} + \frac{\mathbf{y}^3}{6} + \dots\right) \exp\left(-\log n - \frac{\overline{\mathbf{y}^2}}{2} - \frac{\overline{\mathbf{y}^3}}{6} - \dots\right) \\ &= \frac{1}{n} \left(1 + \mathbf{y} + \left(\frac{\mathbf{y}^2}{2} - \frac{\overline{\mathbf{y}^2}}{2}\right) + \left(\frac{\mathbf{y}^3}{6} - \frac{\overline{\mathbf{y}^3}}{6} - \frac{\overline{\mathbf{y}^2} \mathbf{y}}{2}\right) + \dots\right) \end{aligned} \quad (10)$$

Sparse Attention

So what are the applications of these two expansions? First is the one for $\text{logsumexp}(\mathbf{x})$, which is related to the block scoring of Block Sparse Attention like MoBA. We know that Full Attention scores each Token according to $e^{\mathbf{q} \cdot \mathbf{k}}$, so for a certain block $\mathcal{B}$, its score can be reasonably defined as the sum of the scores of each Token, which is equivalent to

$$s_{\mathcal{B}} = \log \sum_{t \in \mathcal{B}} e^{\mathbf{q} \cdot \mathbf{k}_t} \quad (11)$$

Its first-order approximation is $\log |\mathcal{B}| + \mathbf{q} \cdot \bar{\mathbf{k}}$, where $\bar{\mathbf{k}} = \frac{1}{|\mathcal{B}|} \sum_{t \in \mathcal{B}} \mathbf{k}_t$. This exactly corresponds to the empirical practice in MoBA of using AvgPooling as the intra-block Landmark vector. If the first-order approximation feels too coarse, one can consider higher-order corrections. According to Equation (6), the second-order term is

$$\frac{1}{2|\mathcal{B}|} \sum_{t \in \mathcal{B}} (\mathbf{q} \cdot (\mathbf{k}_t - \bar{\mathbf{k}}))^2 = \frac{1}{2} \mathbf{q}^\top \underbrace{\left(\frac{1}{|\mathcal{B}|} \sum_{t \in \mathcal{B}} (\mathbf{k}_t - \bar{\mathbf{k}})(\mathbf{k}_t - \bar{\mathbf{k}})^\top \right)}_{\mathbf{\Sigma}} \mathbf{q} \quad (12)$$

That is, the second-order approximation is $\log |\mathcal{B}| + \mathbf{q} \cdot \bar{\mathbf{k}} + \mathbf{q}^\top \mathbf{\Sigma} \mathbf{q} / 2$, where $\mathbf{\Sigma}$ is exactly the covariance matrix of $\mathbf{k}_t$ within the block. Furthermore, we can consider a diagonal approximation to reduce computational costs. The idea of this higher-order correction is essentially the same as SPLA. Subsequent related work also includes HiLS, which additionally learns a non-equally weighted average central vector to replace $\bar{\mathbf{k}}$.

Linear Attention

As for the application of the $\text{softmax}(\mathbf{x})$ expansion, it is naturally the linearization of Attention. In *Transformer Upgrade Path: 5, Linear Attention as Infinite Dimensions*, we summarized three approaches to linearization, all of which approximate $e^{\mathbf{q} \cdot \mathbf{k}}$. However, after approximating $e^{\mathbf{q} \cdot \mathbf{k}}$, it still needs to be normalized. It might be better to directly approximate the normalized $\text{softmax}(\mathbf{x})$.

Let $\mathbf{x} = (\mathbf{q} \cdot \mathbf{k}_1, \dots, \mathbf{q} \cdot \mathbf{k}_n)$, then we have the first-order approximation:

$$\text{softmax}(\mathbf{x})_i \approx \frac{1}{n} + \mathbf{q} \cdot (\mathbf{k}_i - \bar{\mathbf{k}}) \quad (13)$$

This is exactly the scheme discussed in the paper at the beginning, and it is also very intuitive itself: the average vector $\bar{\mathbf{k}}$ provides a baseline for the relative magnitude of attention. Tokens with a similarity greater than $\bar{\mathbf{k}}$ should have increased attention, otherwise it is decreased. To improve accuracy, Based expands $e^{\mathbf{q} \cdot \mathbf{k}}$ to the second order (which just ensures non-negativity, refer here), but from our perspective, considering the second-order approximation of $\text{softmax}(\mathbf{x})$ is more scientific. The second-order term is

$$\frac{1}{2n} \left[(\mathbf{q} \cdot (\mathbf{k}_i - \bar{\mathbf{k}}))^2 - \frac{1}{n} \sum_{j=1}^n (\mathbf{q} \cdot (\mathbf{k}_j - \bar{\mathbf{k}}))^2 \right] \quad (14)$$

where $(\mathbf{q} \cdot (\mathbf{k}_i - \bar{\mathbf{k}}))^2$ can be transformed into an inner product by introducing an outer product:

$$(\mathbf{q} \cdot (\mathbf{k}_i - \bar{\mathbf{k}}))^2 = \mathbf{q}^\top (\mathbf{k}_i - \bar{\mathbf{k}}) (\mathbf{k}_i - \bar{\mathbf{k}})^\top \mathbf{q} = \left\langle \mathbf{q} \mathbf{q}^\top, (\mathbf{k}_i - \bar{\mathbf{k}}) (\mathbf{k}_i - \bar{\mathbf{k}})^\top \right\rangle_F \quad (15)$$

So, similar to Based, truncating the Softmax in Softmax Attention to the second-order approximation also results in linear attention.

Article Summary

This article derived the Taylor expansions of LogSumExp and Softmax respectively, and discussed two of their potential applications.

When reprinting, please include the address of this article: <https://kexue.fm/archives/11814>

For more detailed reprinting matters, please refer to: Scientific Space FAQ