

# Linearizing Softmax Attention into Gated DeltaNet

Su Jianlin

July 21, 2026

In the article “Taylor Expansion of LogSumExp and Softmax”, we introduced the approximate expansion of LogSumExp and Softmax, and thereby obtained a simple scheme for linearizing Softmax Attention. However, in that article, the resulting linear attention was merely in the form of Vanilla Linear Attention.

In this article, we go a step further and attempt to linearize Softmax Attention into a linear attention variant with the Delta Rule. Surprisingly, the final result is not the basic DeltaNet, but directly yields Gated DeltaNet (GDN).

## Problem Background

Consider the variation pattern of the Attention output with respect to Key and Value, given a fixed Query  $\mathbf{q}$ . Introducing the notation

$$\mathbf{o}_t = \sum_{i=1}^t \alpha_{t,i} \mathbf{v}_i, \quad \alpha_{t,i} = \frac{e^{\mathbf{q} \cdot \mathbf{k}_i}}{Z_t}, \quad Z_t = \sum_{i=1}^t e^{\mathbf{q} \cdot \mathbf{k}_i} \quad (1)$$

Using the Softmax expansion introduced in “Taylor Expansion of LogSumExp and Softmax”, we get  $\alpha_{t,i} \approx \frac{1}{t}(1 + \mathbf{q} \cdot (\mathbf{k}_i - \bar{\mathbf{k}}_t))$ . Substituting this into the above equation yields

$$\mathbf{o}_t \approx \frac{1}{t} \sum_{i=1}^t (1 + \mathbf{q} \cdot (\mathbf{k}_i - \bar{\mathbf{k}}_t)) \mathbf{v}_i = \bar{\mathbf{v}}_t + \left( \frac{1}{t} \sum_{i=1}^t \mathbf{v}_i \mathbf{k}_i^\top - \bar{\mathbf{v}}_t \bar{\mathbf{k}}_t^\top \right) \mathbf{q} \quad (2)$$

where  $\bar{\mathbf{k}}_t, \bar{\mathbf{v}}_t$  are the means of the first  $t$  Key and Value vectors, respectively. It can be seen that the final formula is a linear attention, and it is a variant of the earliest Vanilla Linear Attention (hereafter referred to as “VaLA”).

From “A Brief History of Linear Attention: From Imitation, Innovation to Feedback”, we know that after VaLA, there came the Delta Rule and the corresponding DeltaNet, as well as subsequent extensions like GDN and KDA, which are theoretically stronger linear attention mechanisms than VaLA. So, can we directly linearize Softmax Attention into DeltaNet, GDN, etc., to obtain a better approximation?

## Recursive Form

In fact, Softmax Attention itself can be written in the form of a (non-linear) RNN, which we discussed in “Chapter of Space and Time: Treating Attention as a Quadratic Complexity RNN”. This recursive form is not difficult to understand:

$$\mathbf{o}_t = \sum_{i=1}^t \alpha_{t,i} \mathbf{v}_i = \frac{\sum_{i=1}^t e^{\mathbf{q} \cdot \mathbf{k}_i} \mathbf{v}_i}{Z_t} = \frac{(\sum_{i=1}^{t-1} e^{\mathbf{q} \cdot \mathbf{k}_i} \mathbf{v}_i) + e^{\mathbf{q} \cdot \mathbf{k}_t} \mathbf{v}_t}{Z_t} = \frac{Z_{t-1} \mathbf{o}_{t-1} + e^{\mathbf{q} \cdot \mathbf{k}_t} \mathbf{v}_t}{Z_t} \quad (3)$$

Rearranging gives

$$\mathbf{o}_t = \mathbf{o}_{t-1} + \alpha_{t,t}(\mathbf{v}_t - \mathbf{o}_{t-1}) \quad (4)$$

Sharp-eyed readers may have already noticed that the increment on the right-hand side is in the form of  $\mathbf{v}_t - \mathbf{o}_{t-1}$ , which is “the difference between the new observation and the old prediction”. This already faintly resembles the Delta Rule. Note that so far we have not made any approximations; this equation is completely exact, which implies a natural connection between Softmax Attention and the Delta Rule.

It should also be pointed out that Equation (4) is an RNN with a fixed  $\mathbf{q}$ , but in reality, each position has a different  $\mathbf{q}_1, \mathbf{q}_2, \dots, \mathbf{q}_t$ . These  $t$  different  $\mathbf{q}$ ’s must be substituted in, and they respectively run for  $1, 2, \dots, t$  steps to obtain their own outputs. The total complexity is  $1 + 2 + \dots + t = \mathcal{O}(t^2)$ , which is the meaning of a “quadratic complexity RNN”.

The fundamental reason is that for different  $\mathbf{q}$ ’s, the iteration processes are independent, so each  $\mathbf{q}$  must run a complete round, making the total complexity quadratic. To linearize it, we must find a way to separate the computation of  $\mathbf{q}$ , letting the iteration proceed purely within  $\mathbf{k}, \mathbf{v}$ , and then compute the output using the state variables of the iteration process together with  $\mathbf{q}$ . In this way, the iteration process does not depend on  $\mathbf{q}$ , and only needs to be executed once, reducing the complexity to linear.

## Linear Approximation

Next, we apply the first-order approximation of Softmax to  $\alpha_{t,t}$ , yielding

$$\begin{aligned} \mathbf{o}_t &\approx \mathbf{o}_{t-1} + \frac{1}{t}(1 + \mathbf{q} \cdot (\mathbf{k}_t - \bar{\mathbf{k}}_t))(\mathbf{v}_t - \mathbf{o}_{t-1}) \\ &= \left(1 - \frac{1}{t}\right) \mathbf{o}_{t-1} + \frac{1}{t} \mathbf{v}_t + \frac{1}{t}(\mathbf{v}_t - \mathbf{o}_{t-1})(\mathbf{k}_t - \bar{\mathbf{k}}_t)^\top \mathbf{q} \end{aligned} \quad (5)$$

Some readers might wonder: didn’t we just say that the first-order approximation of Softmax leading to VaLA is not accurate enough? Why are we still using a first-order approximation now? The difference is that previously we needed to approximate  $\alpha_{t,i}$ , where the subscripts cover all  $(t, i)$  pairs, but now we only need to approximate  $\alpha_{t,t}$ , which means we are only approximating the diagonal part. The approximation burden is greatly reduced, and the approximation accuracy subsequently increases.

There is another perspective that can help us better understand this point. First, let's rewrite  $\alpha_{t,i}$  as

$$\alpha_{t,i} = \frac{e^{\mathbf{q} \cdot \mathbf{k}_i}}{Z_t} = \frac{e^{\mathbf{q} \cdot \mathbf{k}_i}}{Z_i} \frac{Z_i}{Z_{i+1}} \cdots \frac{Z_{t-1}}{Z_t} = \alpha_{i,i}(1 - \alpha_{i+1,i+1}) \cdots (1 - \alpha_{t,t}) \quad (6)$$

This equation tells us that for  $t > i$ ,  $\alpha_{t,i}$  can be decomposed into the product of  $\alpha_{i,i}$  and a series of  $1 - \alpha_{j,j}$ . Directly making a first-order approximation for  $\alpha_{t,i}$  versus making separate approximations for  $\alpha_{i,i}$  and  $\alpha_{j,j}$  is equivalent to the difference between  $e^{a+b} \approx 1 + a + b$  and  $e^a e^b \approx (1 + a)(1 + b) = 1 + a + b + ab$ . The latter can introduce a cross-term, achieving higher precision.

## Separating the Iteration

However, even with the approximation, this recursion is still quadratic, and we have not yet achieved the goal of separating out  $\mathbf{q}$ . To get closer to it, we look for a solution of the following form

$$\mathbf{o}_t \approx \mathbf{A}_t \mathbf{q} + \mathbf{b}_t \quad (7)$$

where  $\mathbf{A}_t, \mathbf{b}_t$  are independent of  $\mathbf{q}$ , and we agree on the initial conditions  $\mathbf{A}_0 = \mathbf{0}, \mathbf{b}_0 = \mathbf{0}$ . Of course, the exact solution definitely does not look like this; we are purposefully looking for a solution of the above form to serve as an approximation of the exact solution. Substituting this into the aforementioned recursive form gives

$$\mathbf{A}_t \mathbf{q} + \mathbf{b}_t \approx \left(1 - \frac{1}{t}\right) (\mathbf{A}_{t-1} \mathbf{q} + \mathbf{b}_{t-1}) + \frac{1}{t} \mathbf{v}_t + \frac{1}{t} (\mathbf{v}_t - \mathbf{b}_{t-1} - \mathbf{A}_{t-1} \mathbf{q})(\mathbf{k}_t - \bar{\mathbf{k}}_t)^\top \mathbf{q} \quad (8)$$

An intuitive idea is to separate the iterations of  $\mathbf{A}_t, \mathbf{b}_t$  according to the order of  $\mathbf{q}$ , which would yield recursive formulas for  $\mathbf{A}_t, \mathbf{b}_t$ . However, the left-hand side has at most a linear term in  $\mathbf{q}$ , while the right-hand side features a quadratic term like  $\mathbf{A}_{t-1} \mathbf{q}(\mathbf{k}_t - \bar{\mathbf{k}}_t)^\top \mathbf{q}$ , so the separation cannot be carried out completely. We can only obtain

$$\mathbf{b}_t = \left(1 - \frac{1}{t}\right) \mathbf{b}_{t-1} + \frac{1}{t} \mathbf{v}_t \quad (9)$$

$$\mathbf{A}_t = \left(1 - \frac{1}{t}\right) \mathbf{A}_{t-1} + \frac{1}{t} (\mathbf{v}_t - \mathbf{b}_{t-1} - \mathbf{A}_{t-1} \mathbf{q})(\mathbf{k}_t - \bar{\mathbf{k}}_t)^\top \quad (10)$$

Clearly, the iteration of  $\mathbf{b}_t$  is essentially finding the cumulative average of  $\mathbf{v}_t$ , so we can directly write  $\mathbf{b}_t = \bar{\mathbf{v}}_t$ . The problem is that  $\mathbf{q}$  appears on the right-hand side of Equation (10), meaning the assumption that  $\mathbf{A}_t$  does not depend on  $\mathbf{q}$  fails to hold.

## Simple Discard

A simple approach is to assume that  $\mathbf{q}$  is a small quantity and directly discard it, yielding

$$\mathbf{A}_t = \left(1 - \frac{1}{t}\right) \mathbf{A}_{t-1} + \frac{1}{t} (\mathbf{v}_t - \bar{\mathbf{v}}_{t-1})(\mathbf{k}_t - \bar{\mathbf{k}}_t)^\top \quad (11)$$

This also takes the form of a cumulative average, and solving it gives

$$\mathbf{A}_t = \frac{1}{t} \sum_{i=1}^t (\mathbf{v}_i - \bar{\mathbf{v}}_{i-1})(\mathbf{k}_i - \bar{\mathbf{k}}_i)^\top \Rightarrow \mathbf{o}_t \approx \bar{\mathbf{v}}_t + \left( \frac{1}{t} \sum_{i=1}^t (\mathbf{v}_i - \bar{\mathbf{v}}_{i-1})(\mathbf{k}_i - \bar{\mathbf{k}}_i)^\top \right) \mathbf{q} \quad (12)$$

This looks like a new VaLA variant, but it is actually completely equivalent to Equation (2)! To prove this, we only need to verify

$$t\bar{\mathbf{v}}_t\bar{\mathbf{k}}_t^\top - (t-1)\bar{\mathbf{v}}_{t-1}\bar{\mathbf{k}}_{t-1}^\top = \underbrace{t\bar{\mathbf{v}}_t}_{(t-1)\bar{\mathbf{v}}_{t-1} + \mathbf{v}_t} \bar{\mathbf{k}}_t^\top - \bar{\mathbf{v}}_{t-1} \underbrace{(t-1)\bar{\mathbf{k}}_{t-1}^\top}_{t\bar{\mathbf{k}}_t^\top - \mathbf{k}_t^\top} = -\bar{\mathbf{v}}_{t-1}\bar{\mathbf{k}}_t^\top + \mathbf{v}_t\bar{\mathbf{k}}_t^\top + \bar{\mathbf{v}}_{t-1}\mathbf{k}_t^\top \quad (13)$$

And thus

$$\sum_{i=1}^t \mathbf{v}_i\mathbf{k}_i^\top - t\bar{\mathbf{v}}_t\bar{\mathbf{k}}_t^\top = \sum_{i=1}^t \mathbf{v}_i\mathbf{k}_i^\top - \sum_{i=1}^t (-\bar{\mathbf{v}}_{i-1}\bar{\mathbf{k}}_i^\top + \mathbf{v}_i\bar{\mathbf{k}}_i^\top + \bar{\mathbf{v}}_{i-1}\mathbf{k}_i^\top) = \sum_{i=1}^t (\mathbf{v}_i - \bar{\mathbf{v}}_{i-1})(\mathbf{k}_i - \bar{\mathbf{k}}_i)^\top \quad (14)$$

This transforms Equation (2) identically into Equation (12).

## Error Principle

Another, more accurate approach is to replace  $\mathbf{q}$  with some appropriate quantity  $\mathbf{c}$ . Intuitively, replacing  $\mathbf{q}$  with  $\mathbf{q}_t$  might seem like a natural and good choice, but that is not the case. Let us return to Equation (4); writing  $\mathbf{q}$  explicitly gives

$$\mathbf{o}_t(\mathbf{q}) = \mathbf{o}_{t-1}(\mathbf{q}) + \alpha_{t,t}(\mathbf{q}) \cdot (\mathbf{v}_t - \mathbf{o}_{t-1}(\mathbf{q})) \quad (15)$$

Our goal is to replace the last  $\mathbf{o}_{t-1}(\mathbf{q})$  with some  $\mathbf{o}_{t-1}(\mathbf{c})$ , so that no quadratic term of  $\mathbf{q}$  will appear during the subsequent approximate expansion. The error introduced by this is

$$\alpha_{t,t}(\mathbf{q}) \cdot (\mathbf{o}_{t-1}(\mathbf{q}) - \mathbf{o}_{t-1}(\mathbf{c})) \quad (16)$$

Why might  $\mathbf{c} = \mathbf{q}_t$  not be good? Because it only guarantees the lowest error for the single position where  $\mathbf{q} = \mathbf{q}_t$ , but the state variable at this moment also has to serve the subsequent  $\mathbf{q}_{t+1}, \mathbf{q}_{t+2}, \dots$ . Therefore, the ideal goal is “to minimize the error for any possible  $\mathbf{q}$  that might appear”. This is the “minimum error principle” we use to search for  $\mathbf{c}$ .

This error term has a multiplicative form: if  $\alpha_{t,t}(\mathbf{q})$  is inherently small, then the product will not be large; conversely, if  $\alpha_{t,t}(\mathbf{q})$  is large, then  $\mathbf{o}_{t-1}(\mathbf{q}) - \mathbf{o}_{t-1}(\mathbf{c})$  must be small to bring the error down. From this, we derive a “minimax strategy”: first find the  $\mathbf{q}^*$  that maximizes  $\alpha_{t,t}(\mathbf{q})$ , and then simply set  $\mathbf{c} = \mathbf{q}^*$  to achieve  $\mathbf{o}_{t-1}(\mathbf{q}) - \mathbf{o}_{t-1}(\mathbf{c}) = \mathbf{0}$ . This balances the error on both sides of large and small  $\alpha_{t,t}(\mathbf{q})$ .

## Grand Finale

According to the linear approximation  $\alpha_{t,t} \approx \frac{1}{t}(1 + \mathbf{q} \cdot (\mathbf{k}_t - \bar{\mathbf{k}}_t))$ , the direction of  $\mathbf{q}$  that maximizes it is  $\frac{\mathbf{k}_t - \bar{\mathbf{k}}_t}{\|\mathbf{k}_t - \bar{\mathbf{k}}_t\|}$ . However, since the magnitude can be arbitrarily large,  $\alpha_{t,t}$  effectively has no maximum. To obtain a finite result, we need to constrain the magnitude of  $\mathbf{q}$ . For simplicity, we assume that the magnitudes of  $\mathbf{q}, \mathbf{k}$  are comparable, so we can directly take

$$\mathbf{q}^* = \mathbf{k}_t - \bar{\mathbf{k}}_t \quad (17)$$

This extremely simple form, when substituted into Equation (10), gives

$$\begin{aligned} \mathbf{A}_t &= \left(1 - \frac{1}{t}\right) \mathbf{A}_{t-1} + \frac{1}{t}(\mathbf{v}_t - \bar{\mathbf{v}}_{t-1} - \mathbf{A}_{t-1}(\mathbf{k}_t - \bar{\mathbf{k}}_t))(\mathbf{k}_t - \bar{\mathbf{k}}_t)^\top \\ &= \mathbf{A}_{t-1} \left( \left(1 - \frac{1}{t}\right) \mathbf{I} - \frac{1}{t}(\mathbf{k}_t - \bar{\mathbf{k}}_t)(\mathbf{k}_t - \bar{\mathbf{k}}_t)^\top \right) + \frac{1}{t}(\mathbf{v}_t - \bar{\mathbf{v}}_{t-1})(\mathbf{k}_t - \bar{\mathbf{k}}_t)^\top \end{aligned} \quad (18)$$

This is exactly the appearance of Gated DeltaNet (GDN)! The standard form of GDN is

$$\mathbf{S}_t = \mathbf{S}_{t-1}(\alpha_t(\mathbf{I} - \beta_t \mathbf{k}_t \mathbf{k}_t^\top)) + \beta_t \mathbf{v}_t \mathbf{k}_t^\top \quad (19)$$

The difference lies in whether  $\alpha_t$  multiplies  $\mathbf{I} - \beta_t \mathbf{k}_t \mathbf{k}_t^\top$  or just multiplies  $\mathbf{I}$ , but this difference is not fundamental, and the two can be transformed into each other. At this point, we have completed the approximate transformation from Softmax Attention to Delta Rule-based attention, and our final output is

$$\mathbf{o}_t = \mathbf{A}_t \mathbf{q}_t + \bar{\mathbf{v}}_t \quad (20)$$

## Article Summary

Starting from the concept that “Softmax Attention is a quadratic complexity RNN”, this article first discovered that its recursive increment naturally takes the form of the Delta Rule. Then, by applying a first-order approximation to the diagonal elements of Softmax, and finally finding a substitute for the residual  $\mathbf{q}$  according to the “minimum error principle”, we successfully linearized Softmax Attention into the form of Gated DeltaNet.

(Acknowledgments: This article was completed with the guidance of Kimi K3.)

*When reposting, please include the address of this article: <https://kerue.fm/archives/11823>*

*For more detailed reposting matters, please refer to: “Science Space FAQ”*