

Lab Problem: CUDA Self-Attention

Implement a CUDA program to compute a simplified self-attention operation using the Q, K, and V matrices.

Given Q, K, and V, perform the following steps:

- Transpose K to obtain K^T .
- Compute $Q \times K^T$.
- Apply row-wise Softmax.
- Compute $\text{Softmax} \times V$ to obtain the output.

The Softmax operation is:

$$P_{ij} = \exp(S_{ij} - M_i) / \sum_j \exp(S_{ij} - M_i)$$

where $S = QK^T$ and M_i is the maximum value in row i .

CUDA Requirements:

Implement a CUDA kernel for matrix transpose.

Implement a CUDA kernel for $Q \times K^T$.

Implement a CUDA kernel for row-wise Softmax using a reduction.

Implement a CUDA kernel for $\text{Softmax} \times V$.

.

Input

Define N. Generate the matrices using sequential values for easy verification.

Small test case:

$Q = \begin{bmatrix} 1 & 2 \\ 3 & 4 \end{bmatrix}$

$K = \begin{bmatrix} 5 & 6 \\ 7 & 8 \end{bmatrix}$

$V = \begin{bmatrix} 1 & 2 \\ 3 & 4 \end{bmatrix}$

Output

K^T :

5 7

6 8

QK^T :

17 23

39 53

Softmax:

0.00247 0.99753

0.00000083 0.99999917

Output :

2.995 3.995

3.000 4.000