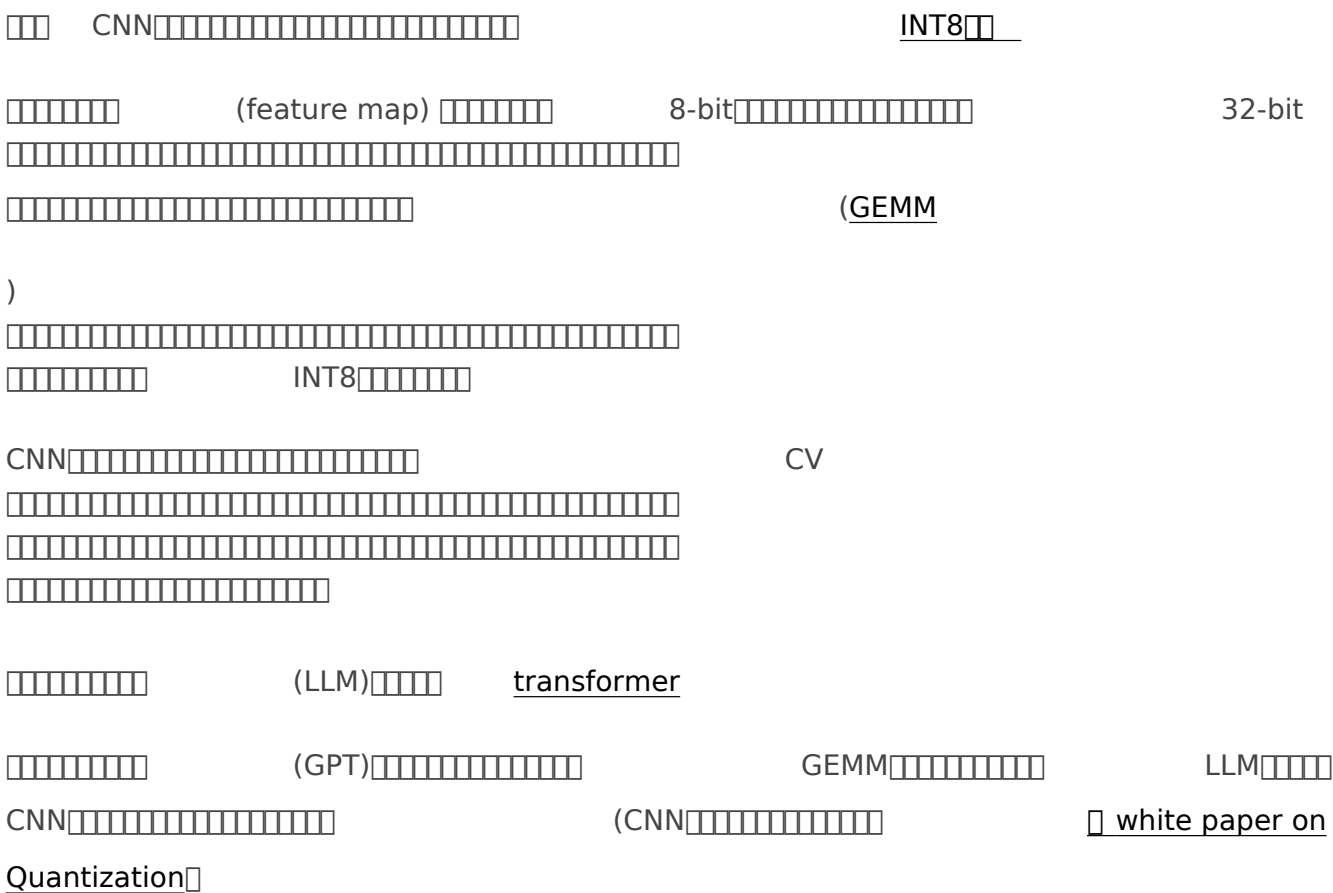
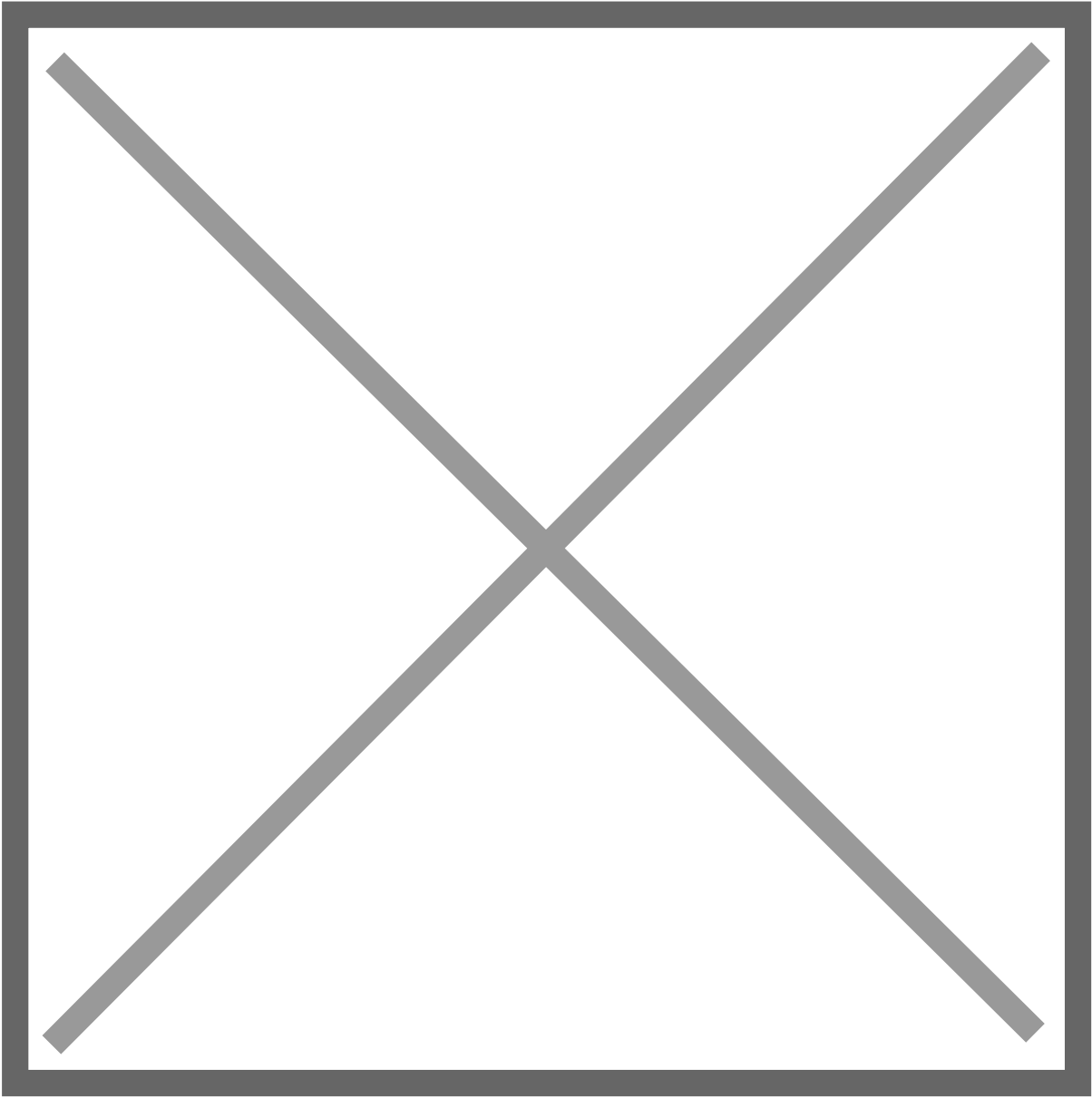


reference



□ CNN □ □ □ □



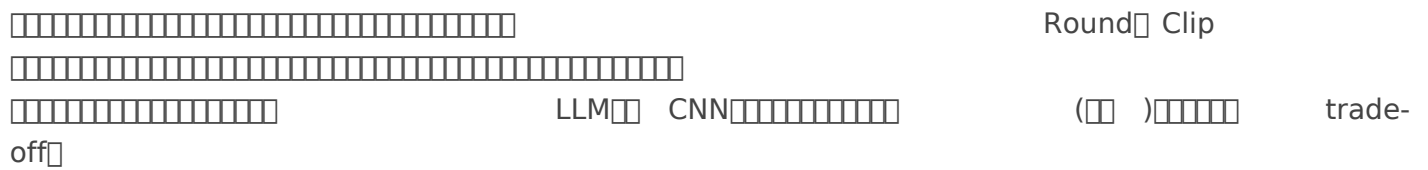
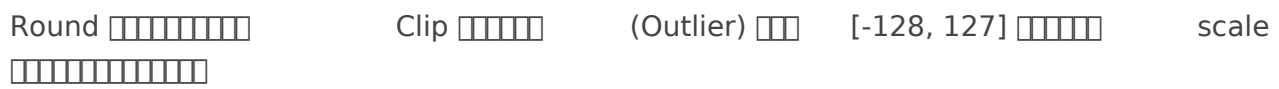


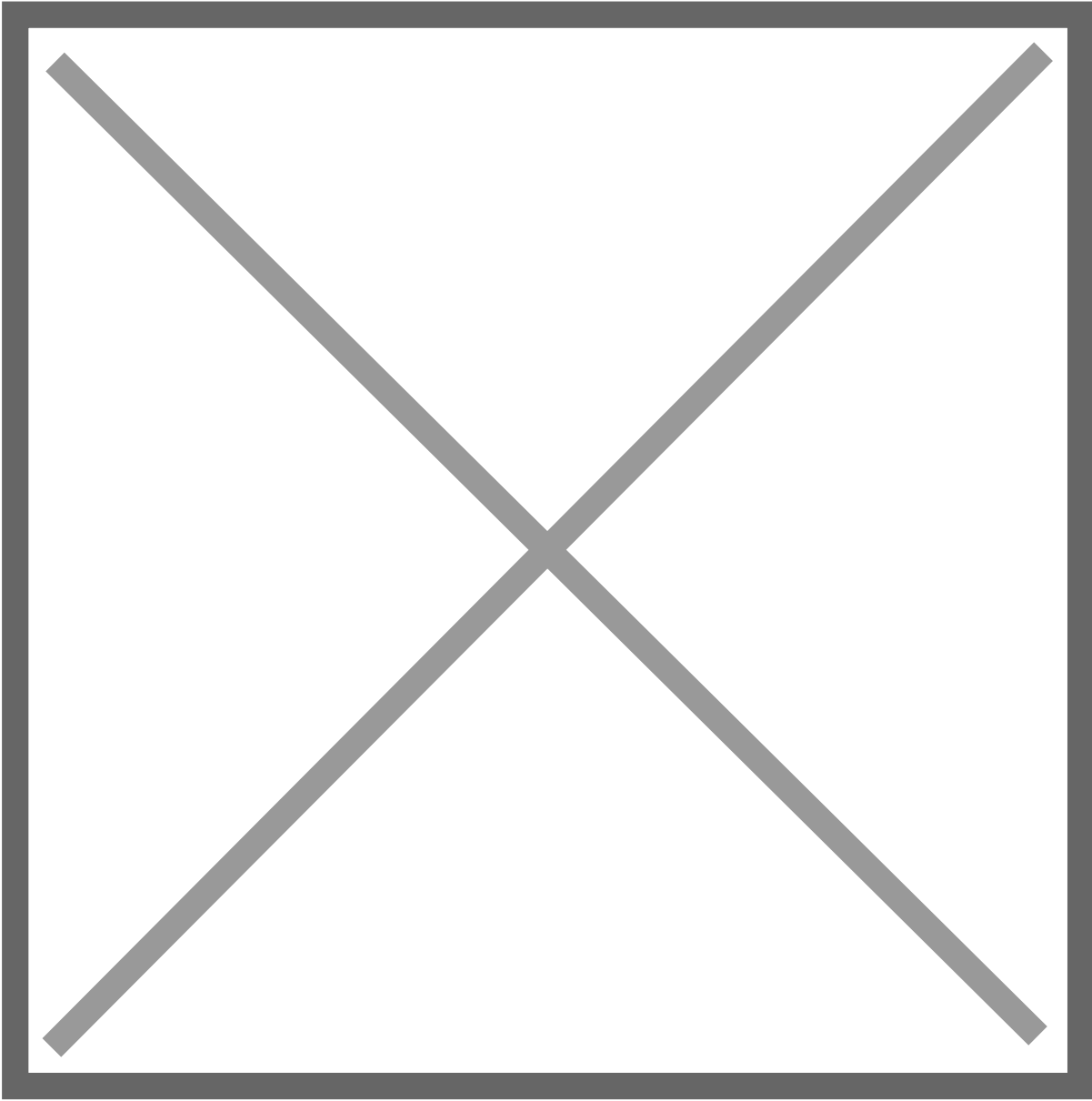
INT8 



[-128, 127]  8bit 







tensor []
scale [] layer [] tensor [] scale [] tensor
[] tensor [] scale []

[] CNN [] tensor [] per-tensor
[] ([] tensor) [] scale
[] per-channel []

[]

[] LLM [] LLM [] per-vector [] per-group [] vector
[]



“ ”

float32 □ float16 □ bfloat16

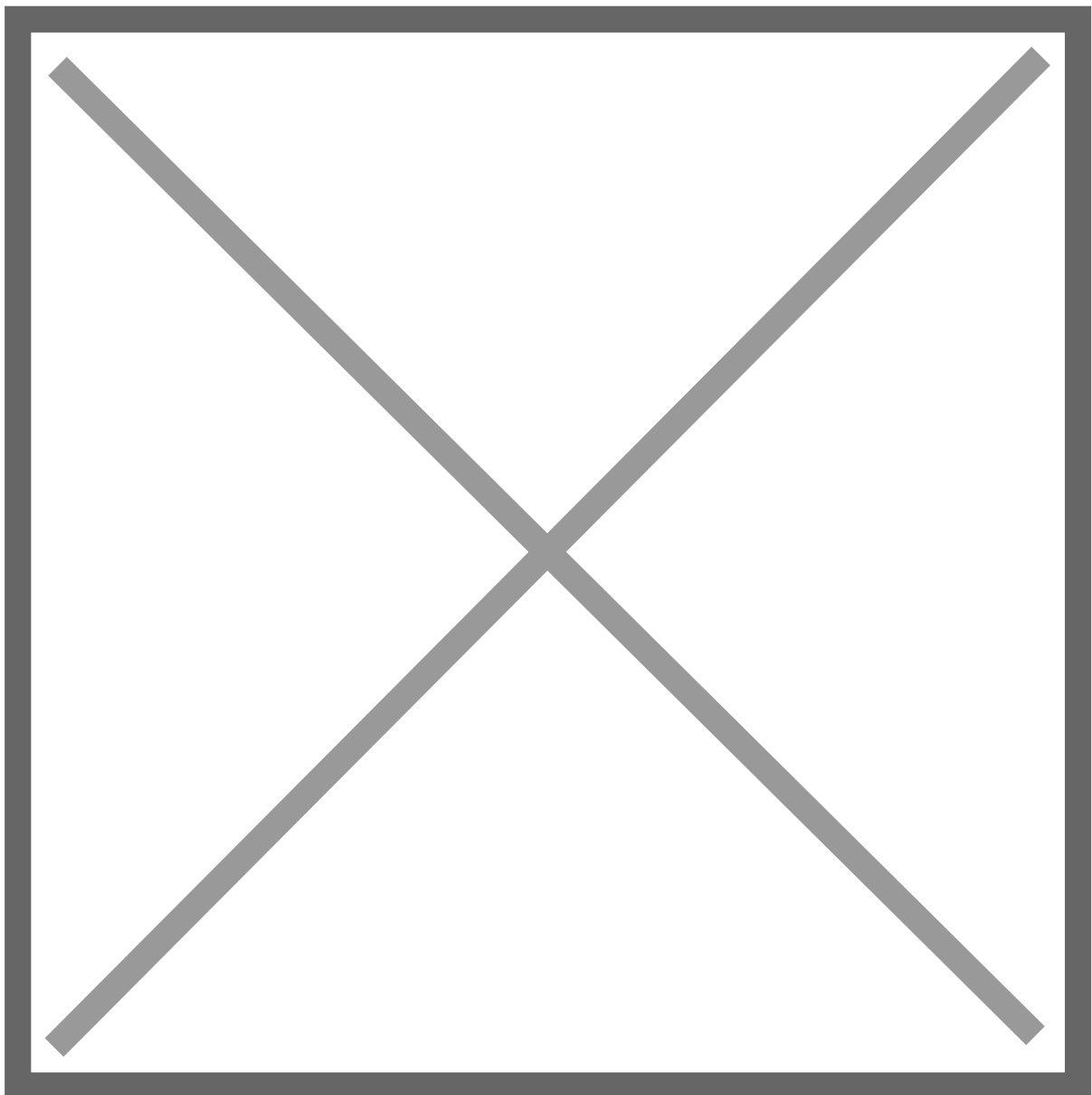
INT8□□□□

INT4

--	--	--	--	--	--	--	--

[illegible]

-  sign  s  1 bit 
-  exponent  k  8 bits  2 
-  fraction  M  23 bits 



●Float32 (FP32)

IEEE 32

--	--	--	--	--	--	--	--

8

--	--	--	--

23

--	--	--	--

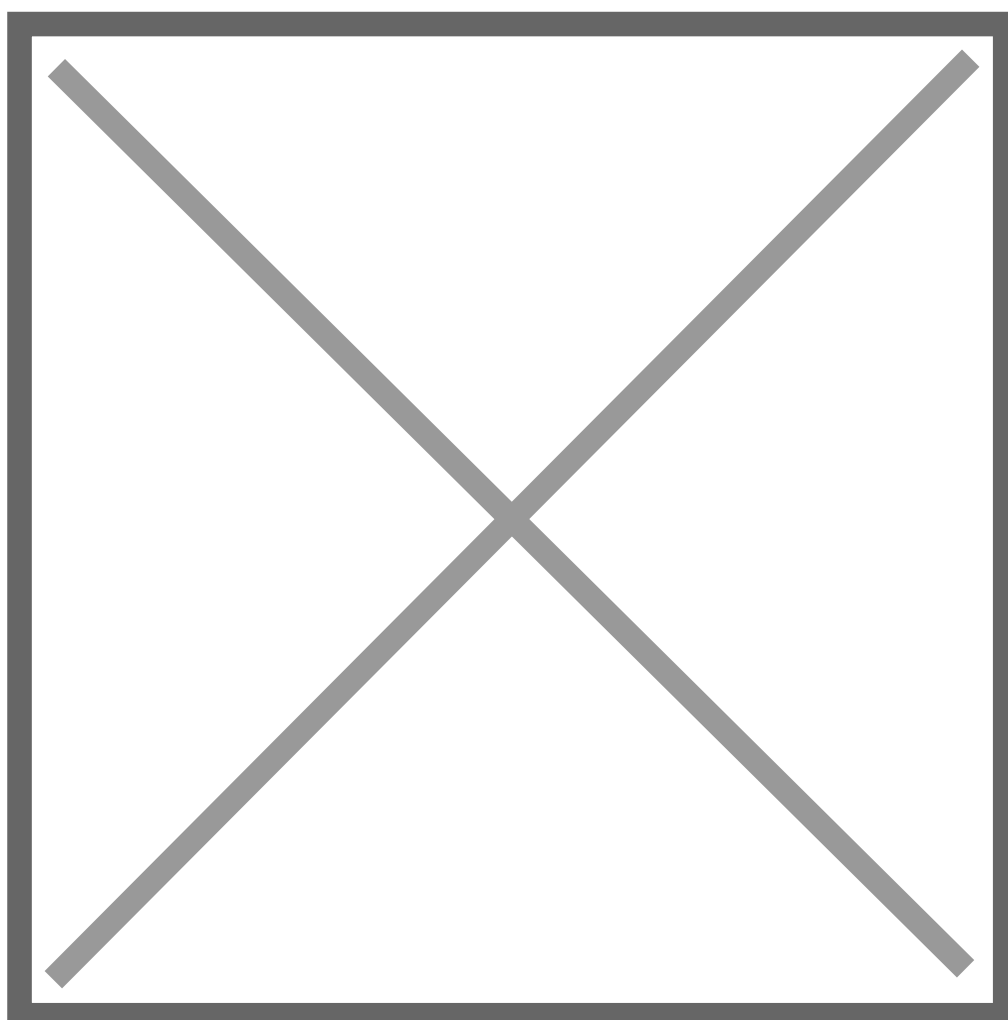
1

FP32

--	--	--	--

● **Unsigned Char (INT4 * 2)**   Unsigned Char  INT4  3

  1  [-8, 7] 



CNN vs. LLM

LLM[] CNN[]

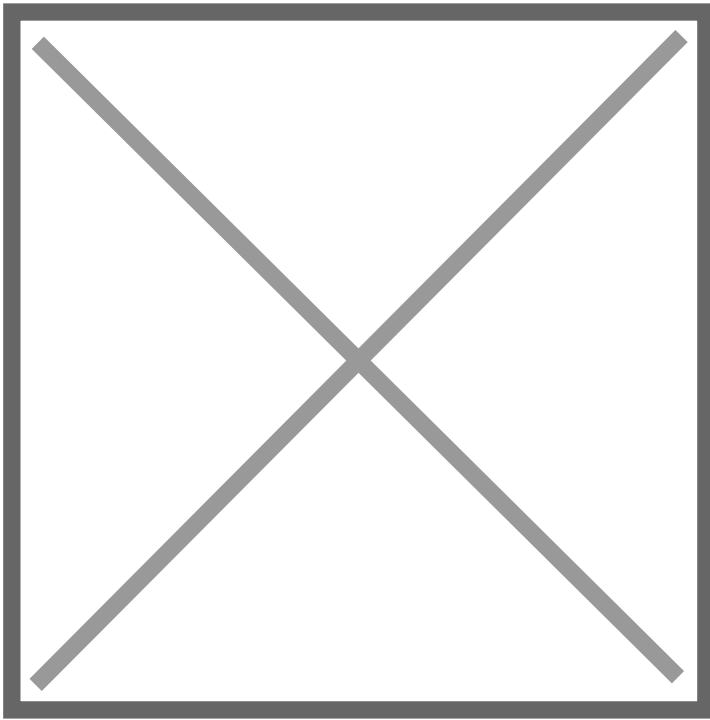
[] [] CNN[] per-channel[] per-layer[] kernel[] LLM
[] [] vector[] (tensor[])
[] A*B=C[]
per-tensor[] A [] B [] INT
[] per-vector[] LLM[] per-vector[] LLM
[] X[] W
[] FP16[] k[] k
[] 2[] k[] 256[] 128[] per-group[]

[]

[] **INT4** CNN[] INT8[] LLM
[] CNN[] INT4
[] 13B[] scale
[] scale[] fp16[] fp32[] scale[] INT4
[]

[] [] CNN[] loss
[] per-channel
[] scale
[] LLM[] head
[] tensor[] hidden dim
[] [50, 100]
[] [0.5, 1.2], [] LLM[]
attention[] softmax[] Attention Is Off By One[] tensor[]
per-tensor[] LLM
[]

[] [] CNN[] LLM[] cuda[] INT
[] Core
[]
[]
[] []



LLM

LLM

LLM.int8():

CNN PTQ (KL)

bit

CNN

LLM tensor

6B

layer

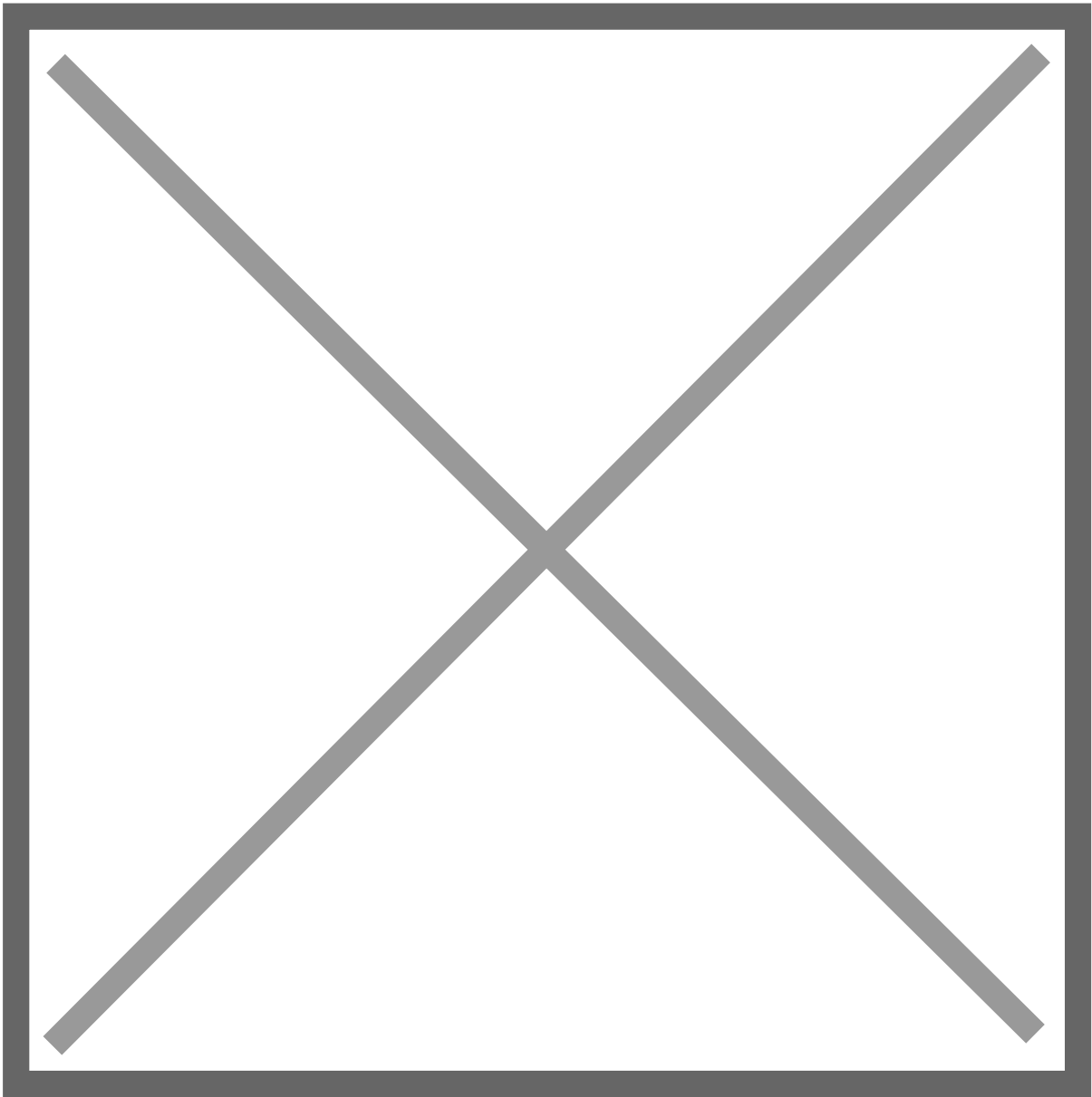
sequence length

quantization

hidden size

sequence token

paper

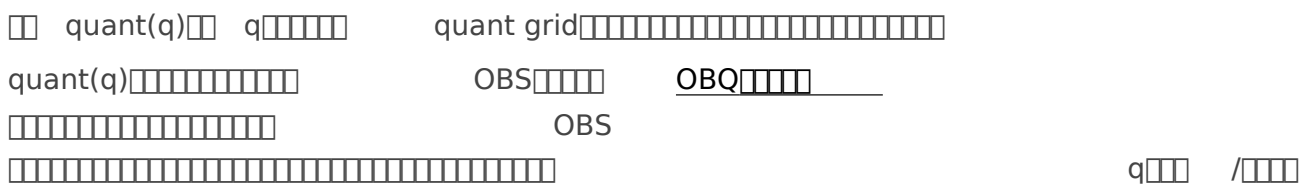


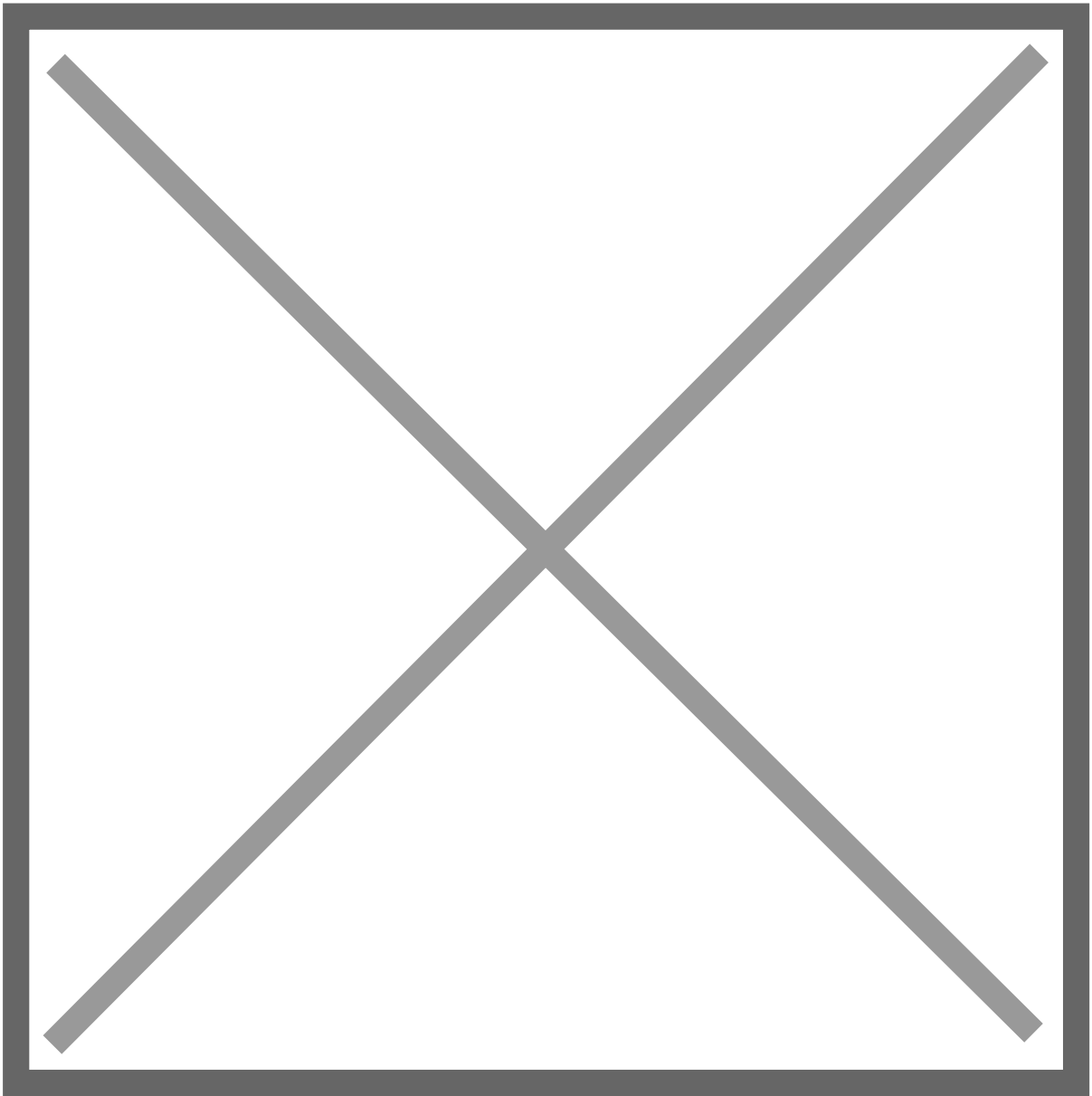
X[] activations[] sequence length[] hidden-size
[]
[]

[] X W[] int8
[] fp16
[] (W[]
token[])
element-wise[] concat[]

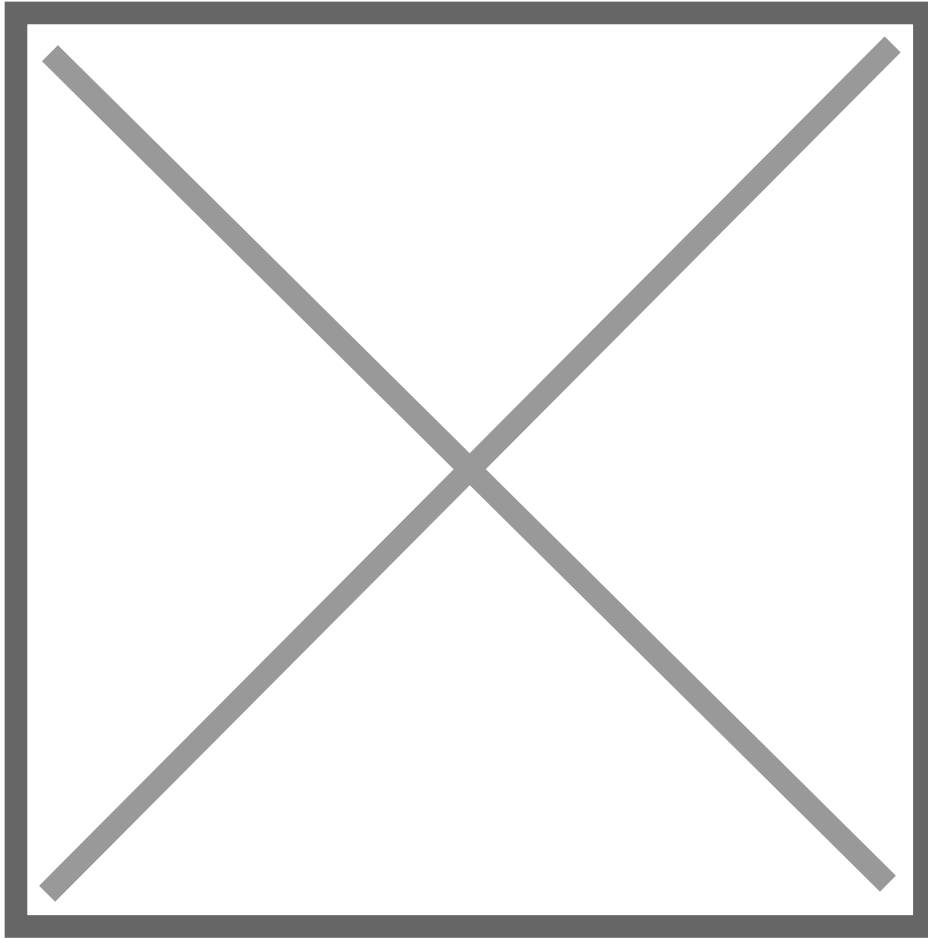
LLM.INT8()[] gemm[] X[] W
[] threshold[] 6
[] 6[]

LLM.int8()[] BLOOM-
176B [] FP16 [] 15% [] 23% [] ([] T5-3B [] T5-11B)

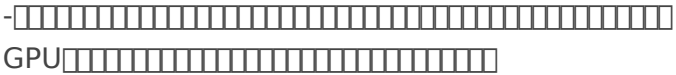
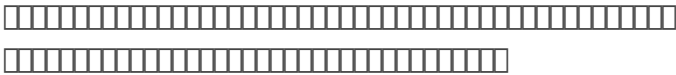




GPTQ ☐ OBS ☐



-



SmoothQuant: W8A8 scale

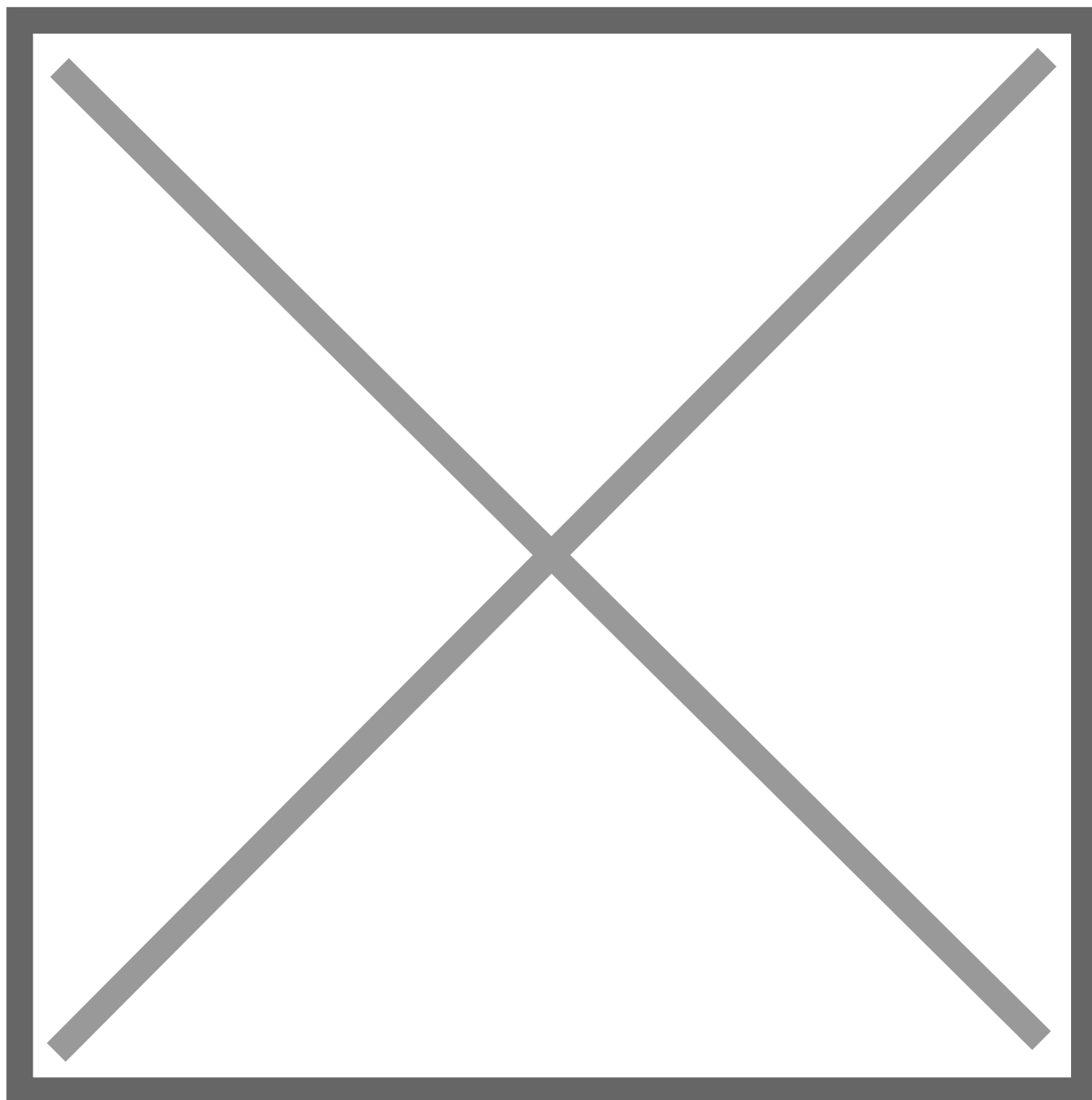
 LLM  channel  activation outlier  (e.g. GLM-130B 
30%  int8  -128  127  outlier  -128  127

[-

4bit

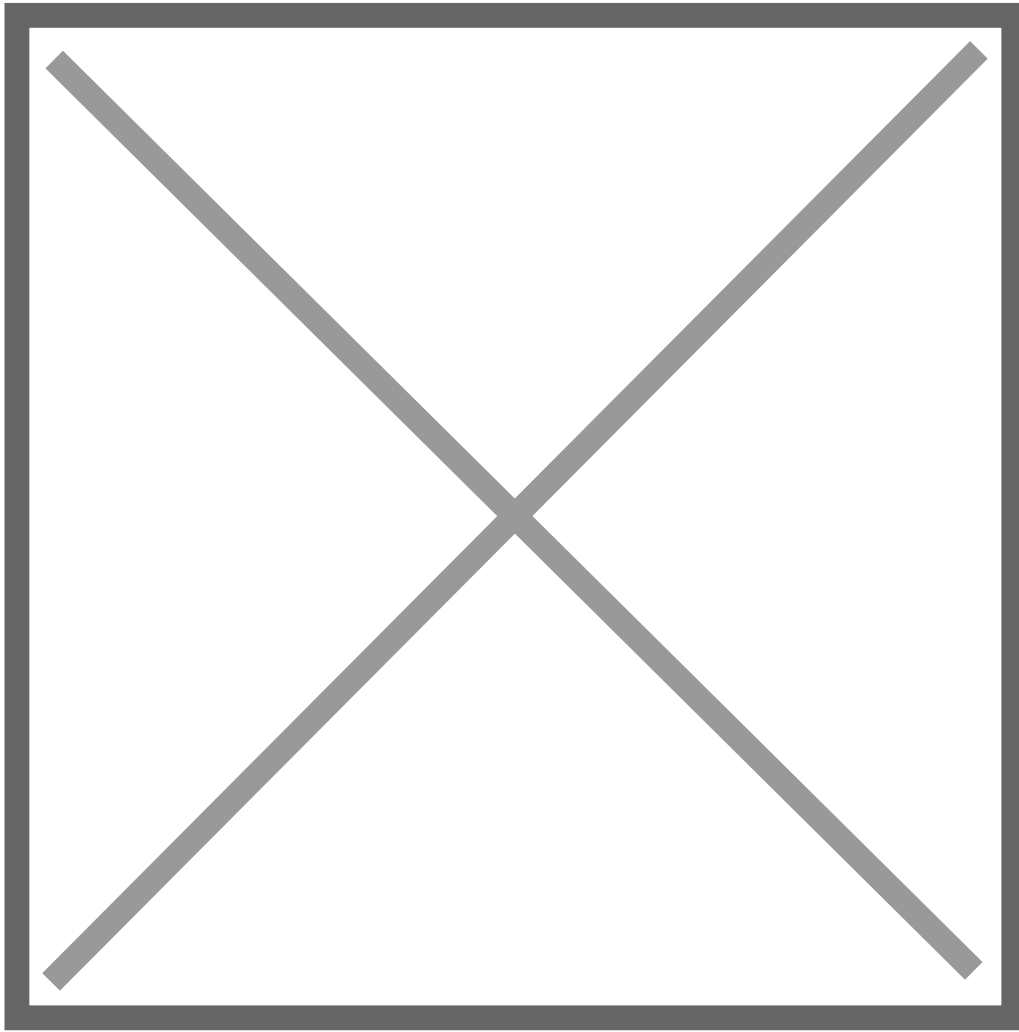
--	--	--	--	--	--	--	--

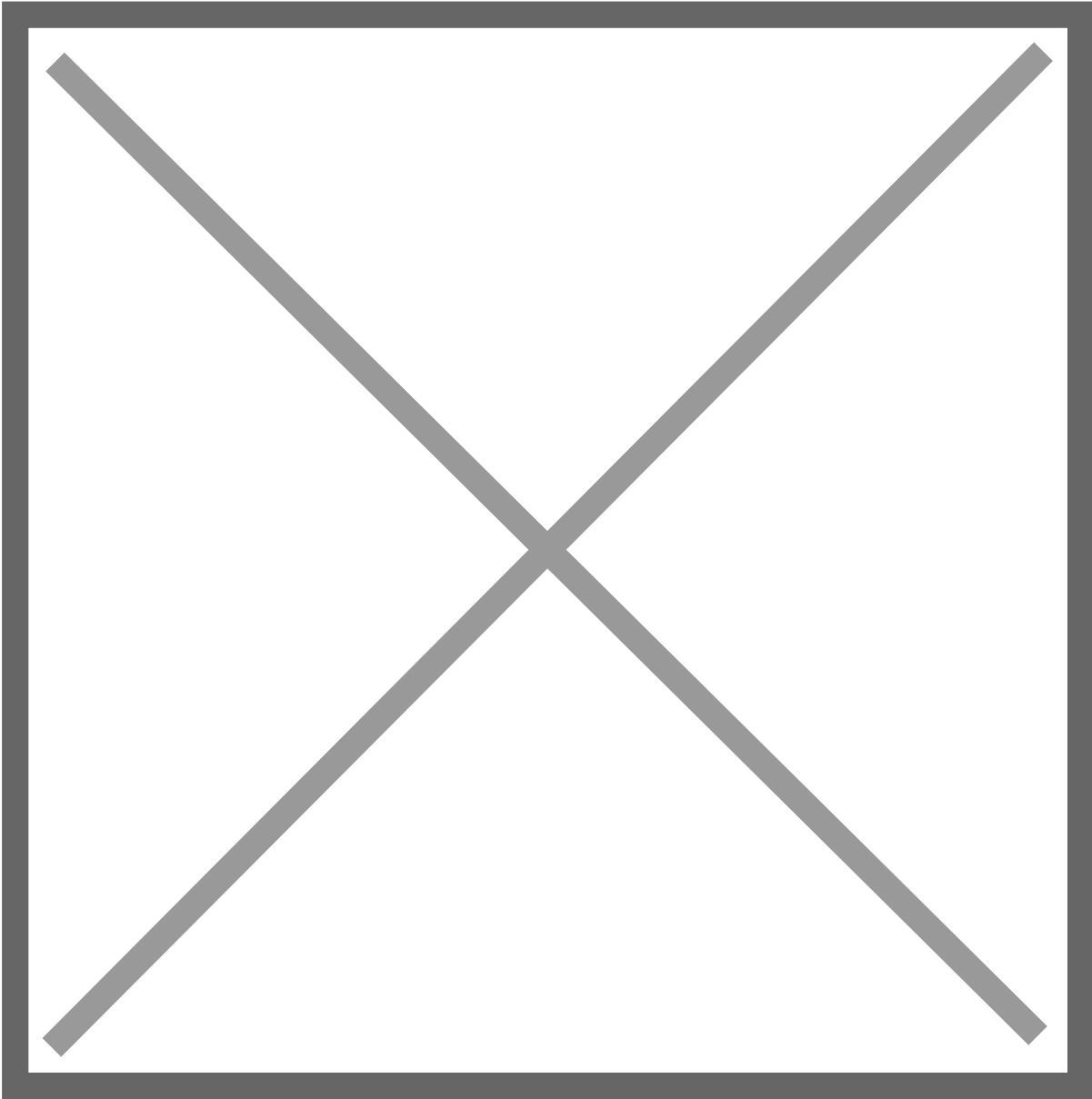
--	--

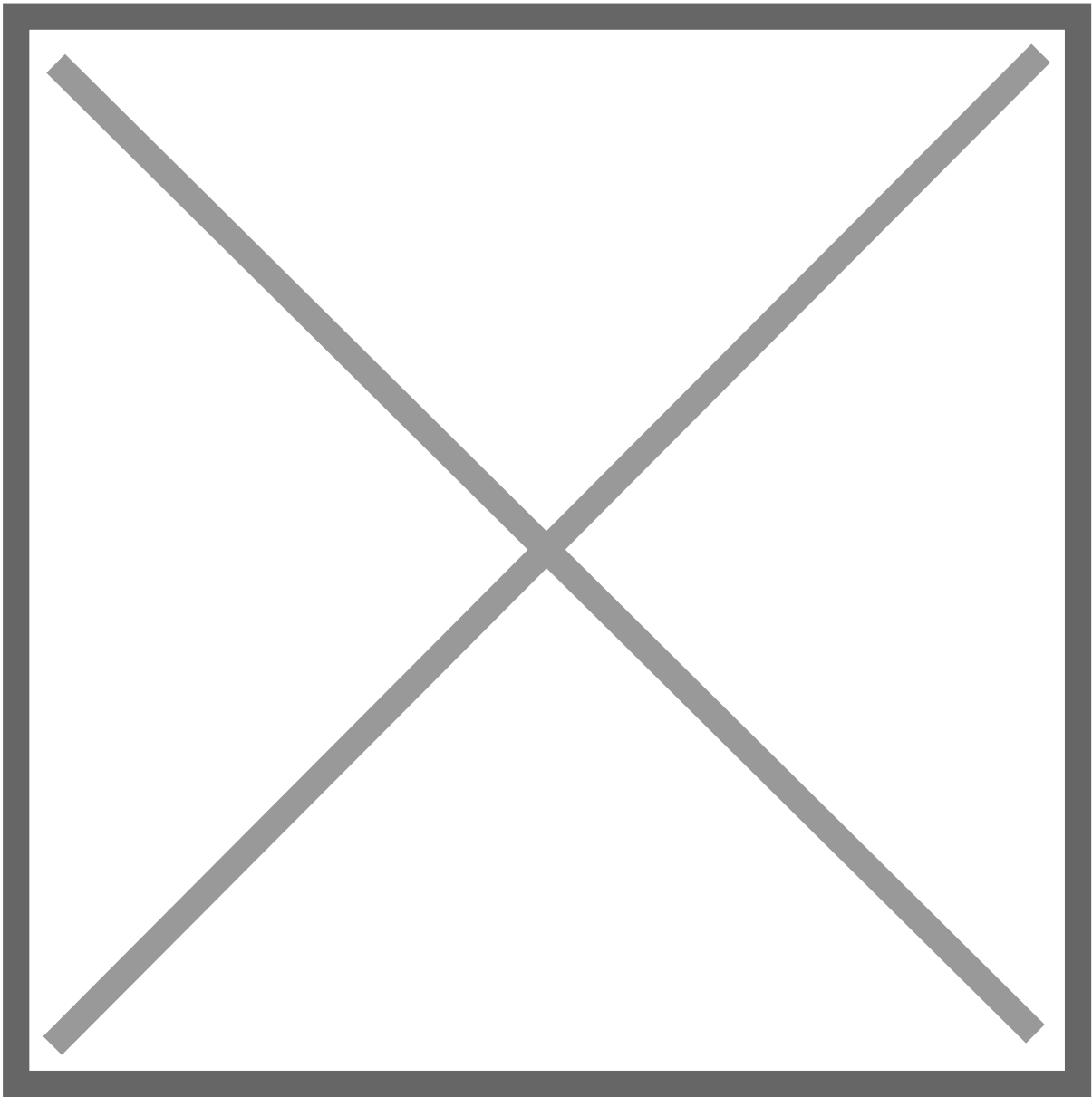
weight 

fp32

S ☐ ☐ ☐ ☐

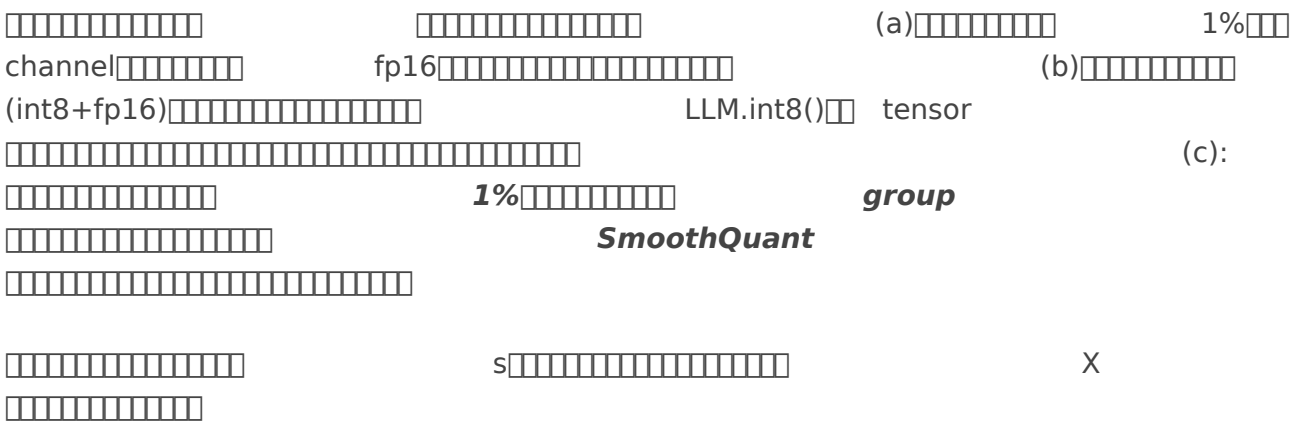


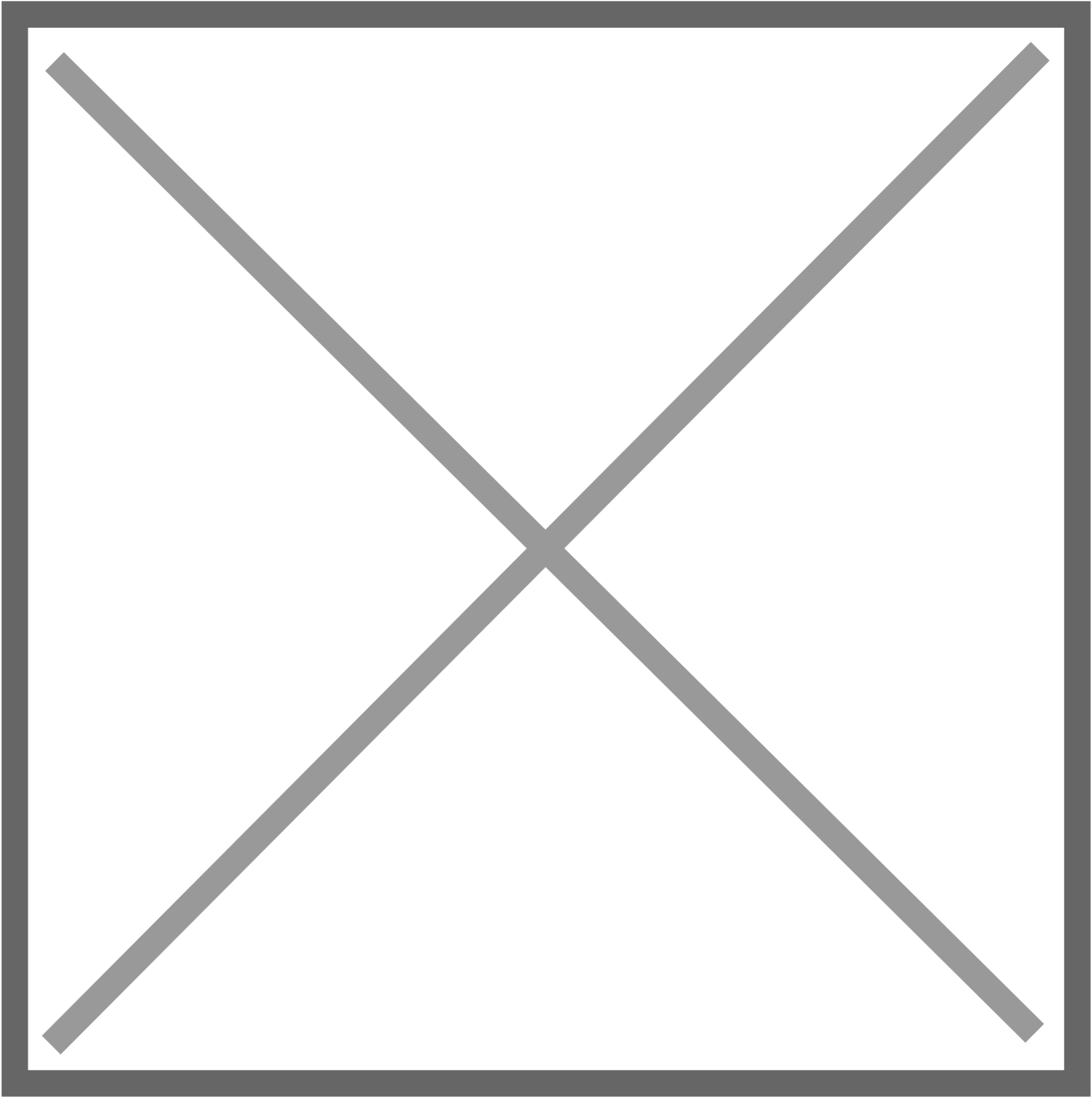




LLM weight-only SmoothQuant LLM
SmoothQuant W8A8
INT4 bit
SmoothQuant LLM
weight-only
SmoothQuant TensorRT-LLM TensorRT LLM
GPT2 SmoothQuant

AWQ:





□□□□□□□

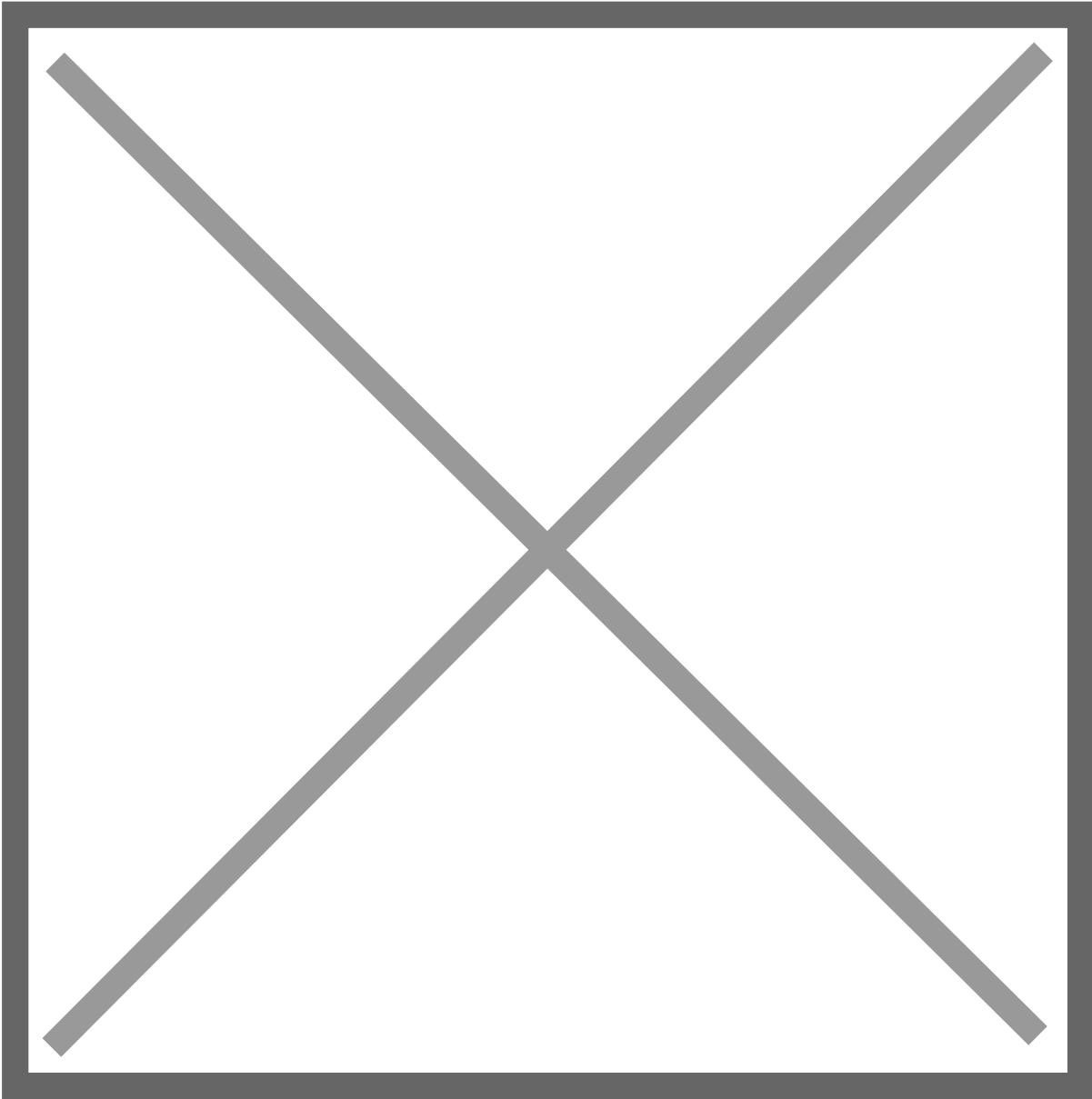
s□□□□□□□□

The diagram illustrates various quantization and pruning techniques for neural network layers. It shows different data flows and operations:

- Top Section:** Shows operations involving $sX = \text{meanc_out}|X|$, $sW = \text{meanc_out}|W|$, sX , group , weight , α , β , and $[0, 1]$. It also shows a sW vector and a group vector.
- Middle Section:** Shows a kernel vector and a weight vector.
- Bottom Section:** Shows operations involving smooth quant , AWQ , weight-only , SQ , and W8A8 . It also shows a kernel vector.

Llama. CPP ☐ K-quant ☐

LLaMA.cpp ██████ Georgi Gerganov █ Meta █ LLaMA █████ C/C++
██████████ tensor█ GGML████████ llama.cpp █████ GPU
████████ Meta █ LLaMA ████████████████████████████████ cuBLAS█████
6█ 15█ PR█ llama.cpp█████ CUDA████████ cuda kernel████ llama█ GPU
██████



llama.cpp██████████████████ INT4████ K-quant██████████████████ Q
████████ x██████ x bit██████████ y █ 0██████████████████ y█ 1
██████████ scale████ zero point█ y█ K████ K-quant████ K
████████████████████ Small, Meduim█ Large████ Q4_0█ Q8_0██████ INT4█
INT8████

QLora

LoraQLoraQLora

lora

ABWLoRAAB

QLora

Block-wise Quantizationblock)scalescaleblock

Quantile Quantization

Quantile Quantization. Quantile

block-wise quantizationblock)scalescale

float32, block6432/64 =

0.5bit4bit0.5bit12.5%

scaleoutliner

256block8bitc8/64 + 32/256

= 0.127 bit

QLoraLLM



- LLMweight-only
- LLMCNNper-vector + scale
- fp16INT

Llama.cpp [] [] [] []

- LLM[] INT4[] + []

INT8[]

INT4

Revision #3

Created 2 March 2025 12:35:11 by Mark

Updated 4 March 2025 01:54:19 by Mark