

Model Compression

Model compression techniques are mainly classified into quantization (Zafrir et al.,2019), pruning (Hoeffler et al., 2021; Vadera and Ameen, 2022), Neural Architecture Search(NAS) (Sun et al., 2019), and knowledge distillation

- [Model Quantization](#)
- [Neural Arch Search](#)
- [Knowledge Distillation](#)
- [Pruning](#)
- [reference](#)

Model Quantization

Quantization Granularity

Quantization is a magic spell to reduce the memory footprint of a model. But often quantization leads to a drop in the accuracy of the model. This is where the Granularity of quantization comes into the picture. Selecting the right granularity helps in maximizing the quantization without much drop in the accuracy performance

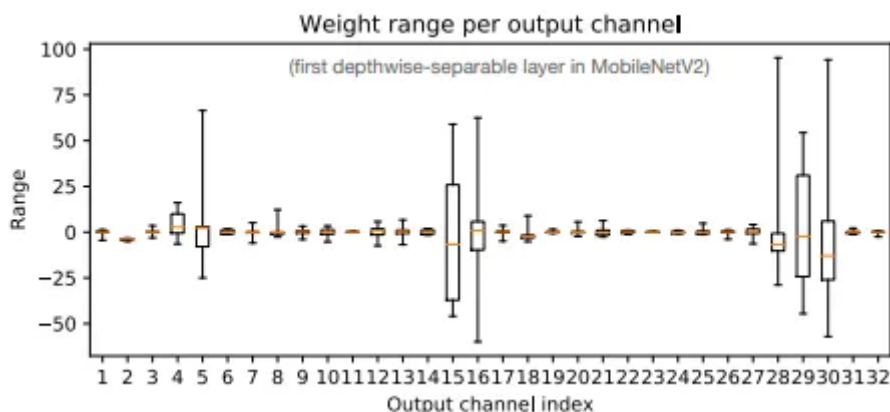
Per-Tensor Quantization

In per-tensor quantization, the same quantization parameters are applied to all elements within a tensor. Applying the same parameters across tensors will cause a drop in accuracy.

Per-Channel Quantization

In per-channel quantization, different quantization parameters are applied to each channel of a tensor independently. This often leads to a lower error while quantizing compared to per tensor quantization.

Per-channel quantization captures variations in different channels more accurately. This usually helps in CNN models where the range of weights varies over different channels.



Mix-precision Quantization

different layer using different precision.

Quantization Method

Quantization Method	Accuracy	Model Size Reduction	Inference Speed	Ease of Implementation
GGUF	High (adaptive quantization of parameter groups)	Significant (grouped quantization)	High (optimized through grouped quantization)	Moderate (requires parameter grouping)
AWQ	High (asymmetric quantization for accuracy)	Moderate	Moderate to High (asymmetric approach)	Moderate (asymmetric scaling)
GPTQ	Moderate (post-training fine-tuning)	Moderate	Moderate (benefits from post-quantization)	Easy (applied post-training)
GGML	High (mixed-precision within groups)	Significant (effective group-wise precision)	High (mixed-precision within groups)	Complex (group-wise mixed-precision)
PTQ	Moderate (quick and easy implementation)	High (straightforward implementation)	High (quick deployment)	Easy (quick application)
QAT	Very High (trained with quantization in mind)	Moderate to High (requires complex training)	High (optimized for quantization during training)	Complex (requires training with quantization)
Dynamic Quantization	Moderate (applied during inference)	High (applied at inference time)	Moderate to High (quick deployment)	Easy (inference-based)
Mixed-Precision	High (selectively applied different precisions)	Moderate to High (selectively applied)	High (selective precision boosts performance)	Complex (selective precision assignment)

Neural Arch Search

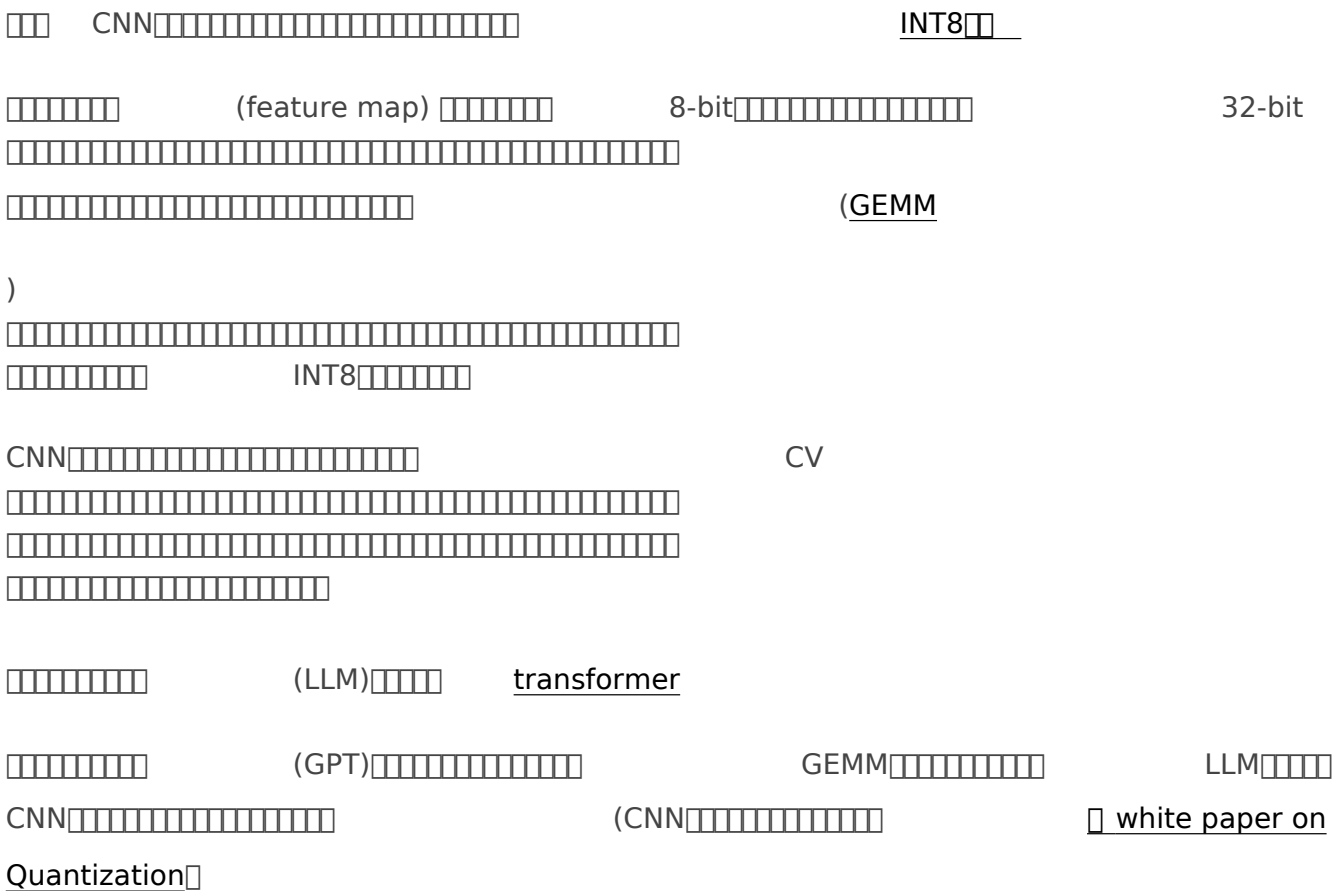
Knowledge Distillation

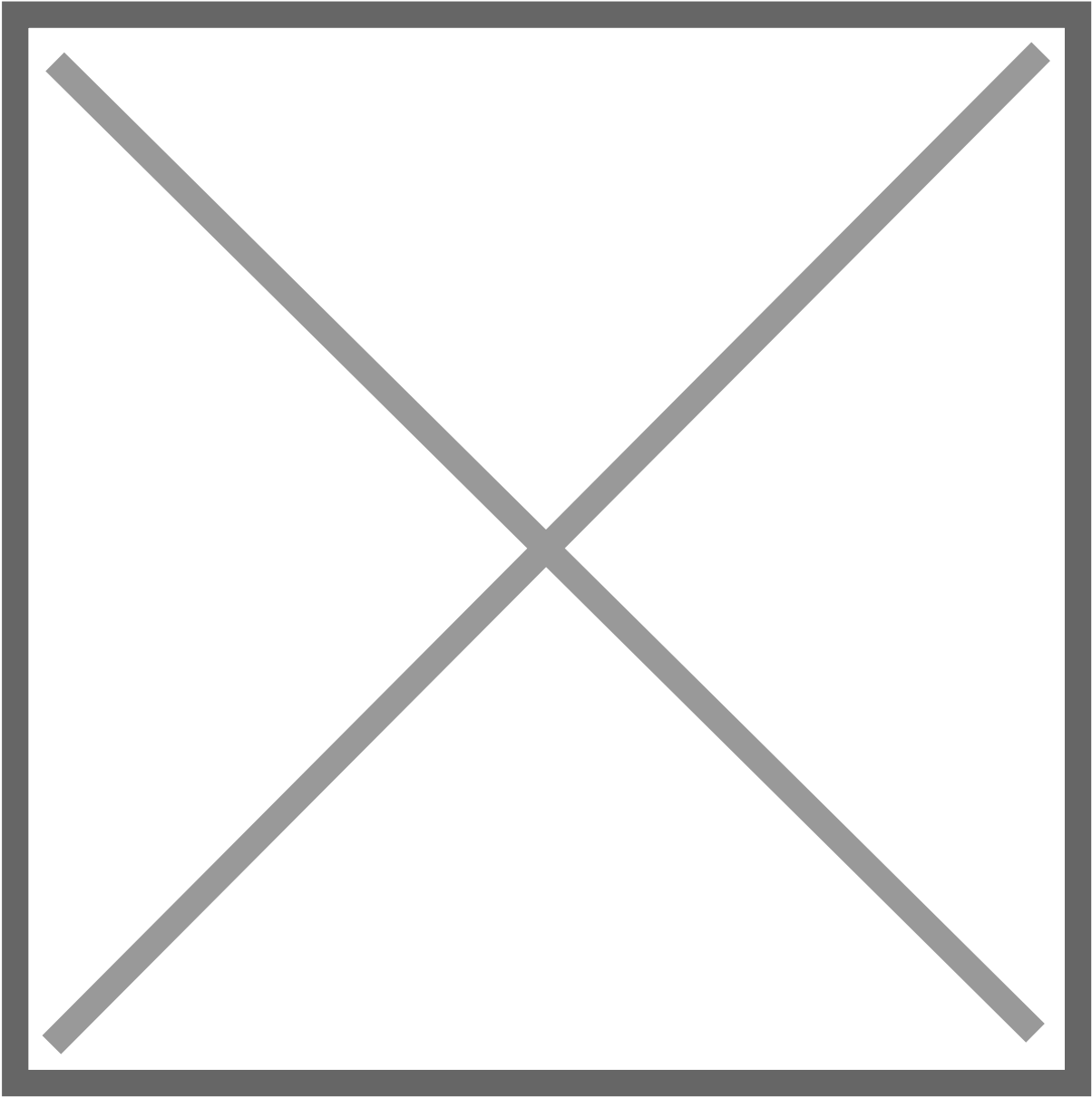
Pruning

reference



□ CNN □ □ □ □



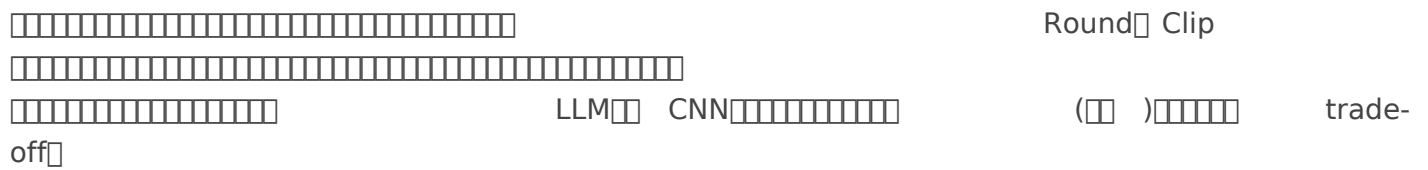
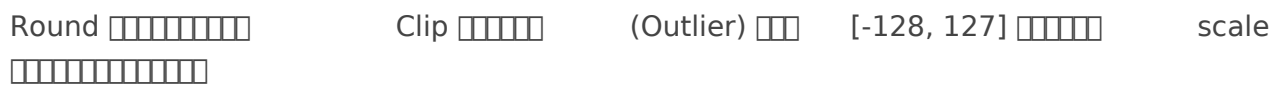


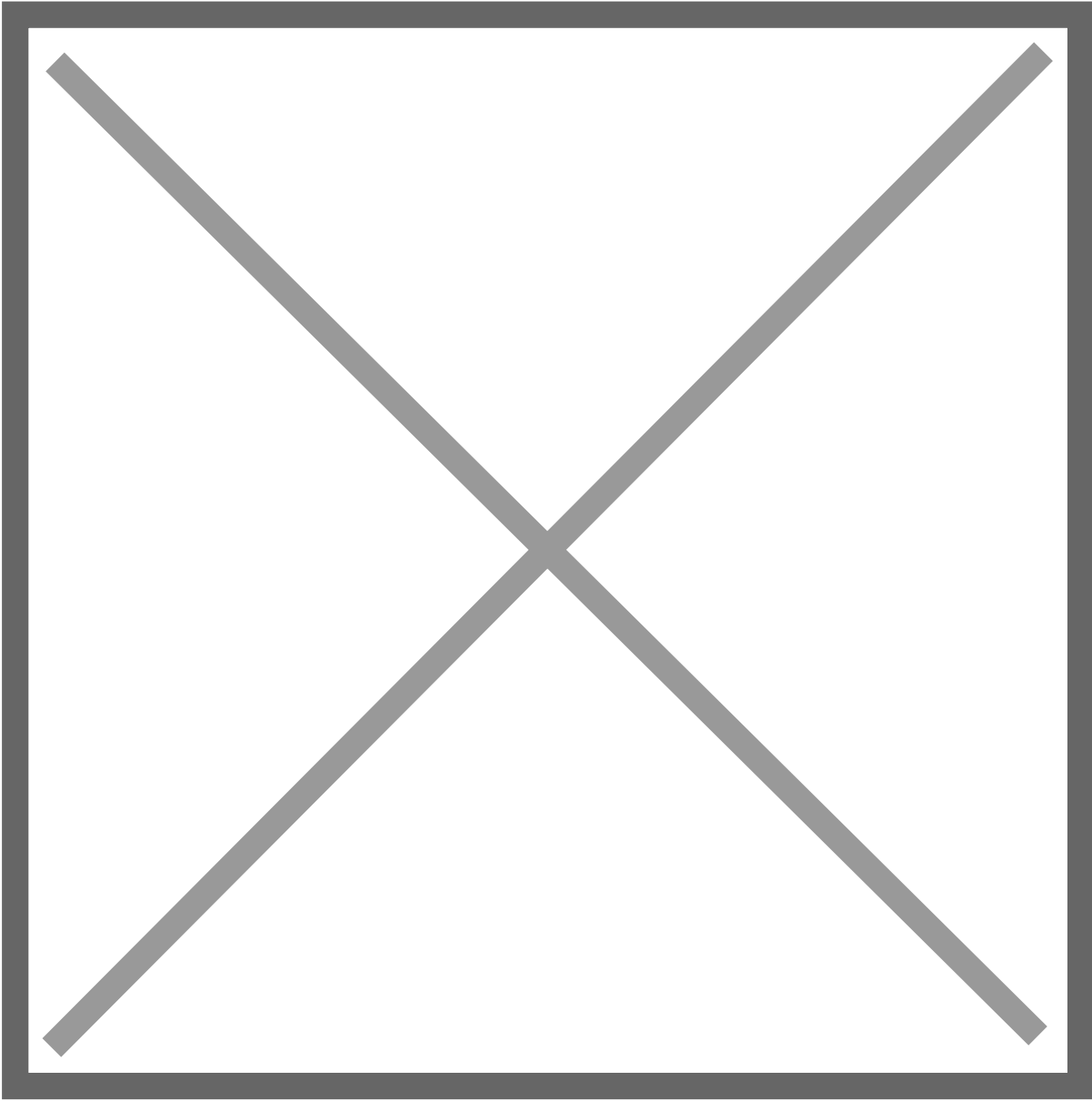
INT8 



[-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

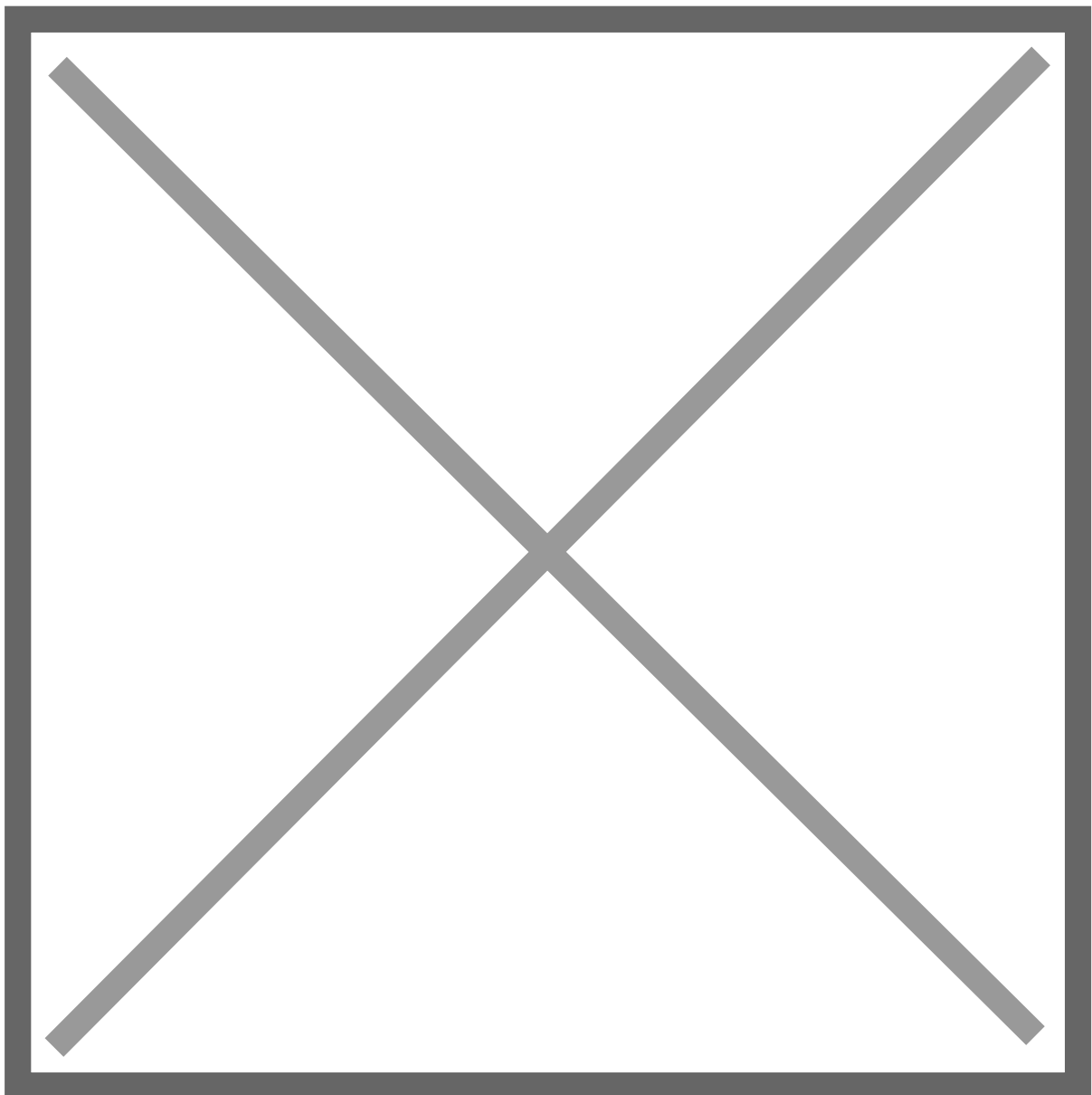
INT8□□□□

INT4

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

LLM

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



●Float32 (FP32)

IEEE 32

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

8 ☐ ☐ ☐ ☐

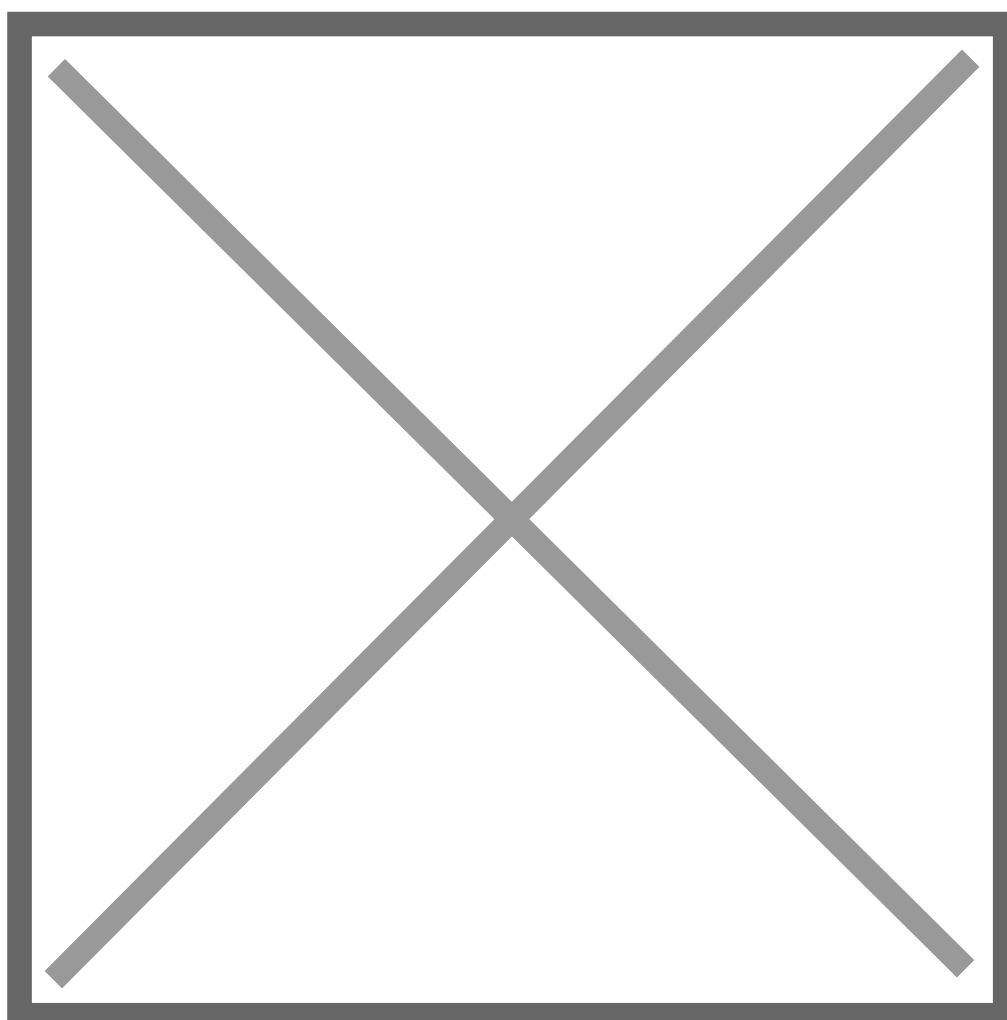
23 ☐ ☐ ☐ ☐

1

FP32

--	--	--	--

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



CNN vs. LLM

LLM[] CNN[]

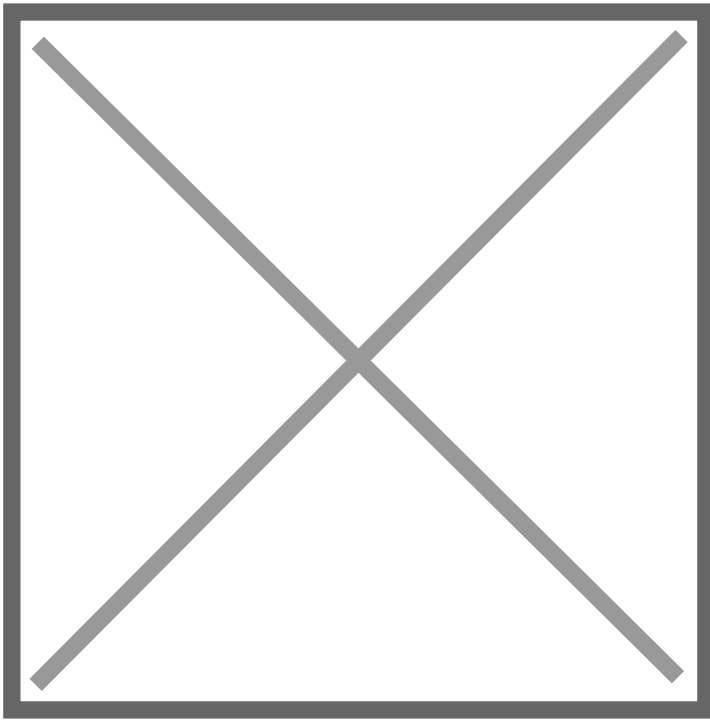
[] [] CNN[] per-channel[] per-layer[] kernel[] LLM
[] [] vector[] (tensor[])
[] A*B=C[]
per-tensor[] A [] B [] INT
[] per-vector[] LLM[] per-vector[] LLM
[] X[] W
[] FP16[] k[] k
[] 2[] k[] 256[] 128[] per-group[]

[]

[] **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

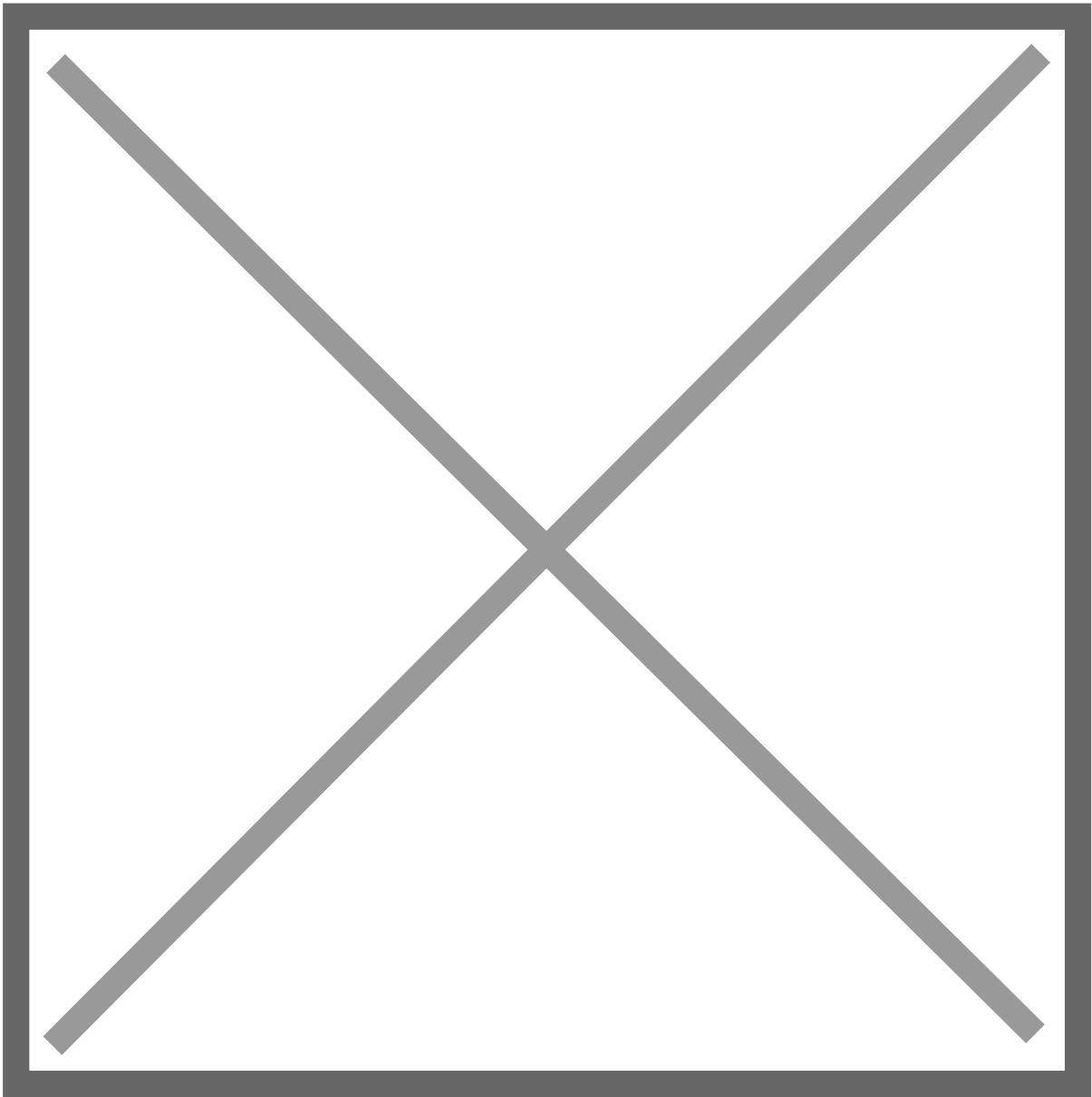
sequence length

quantization

hidden size

sequence token

paper

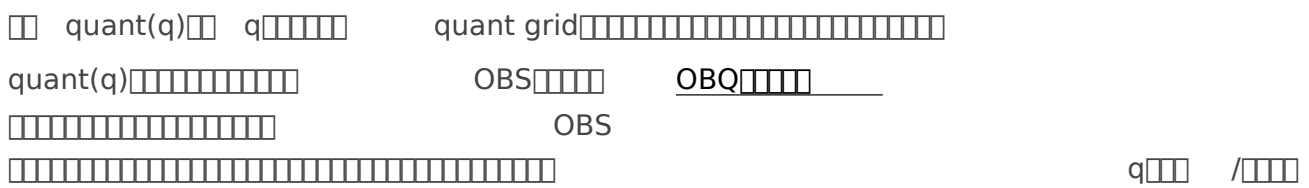


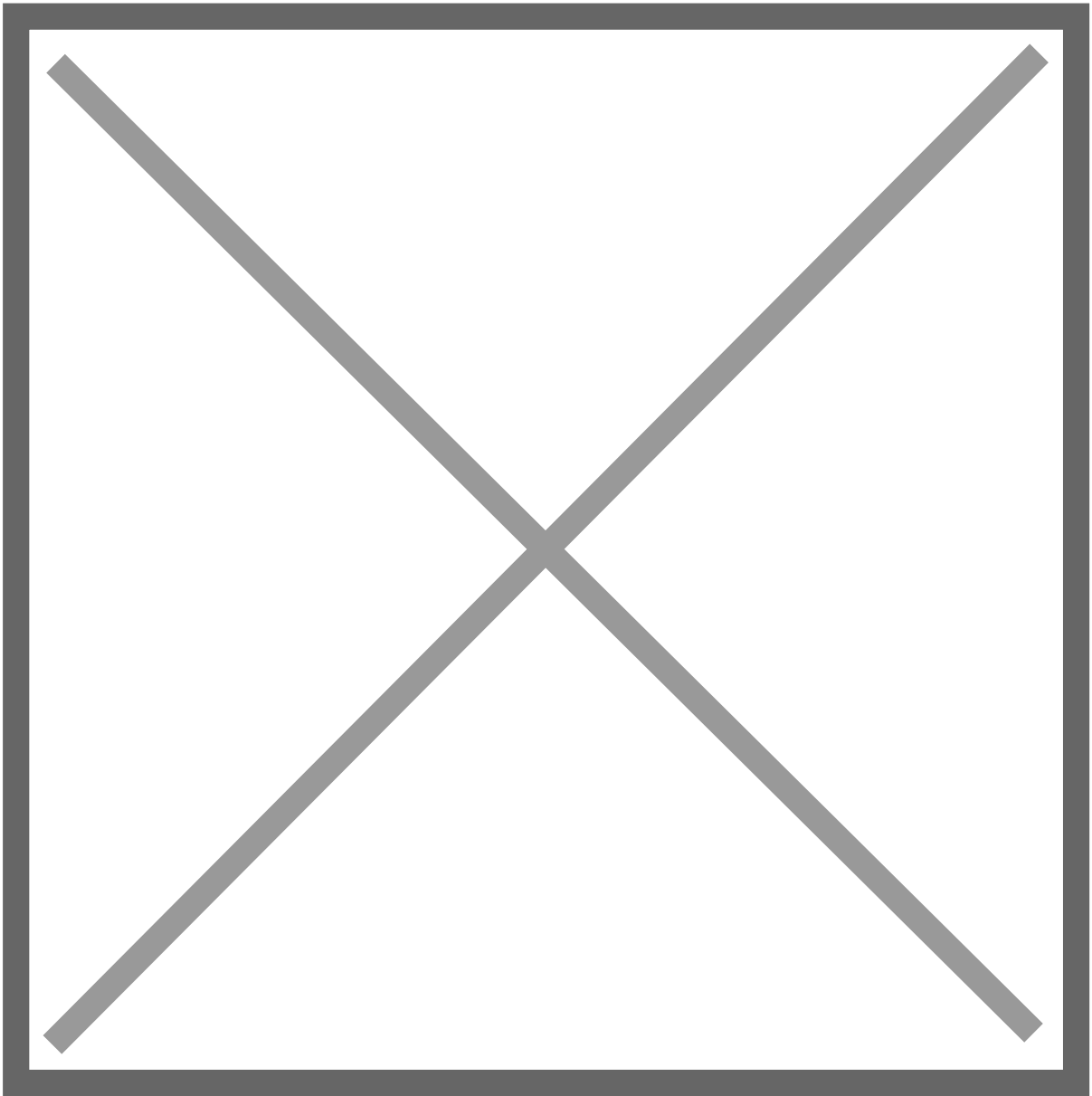
X[] activations[] sequence length[] hidden-size
[]
[]

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

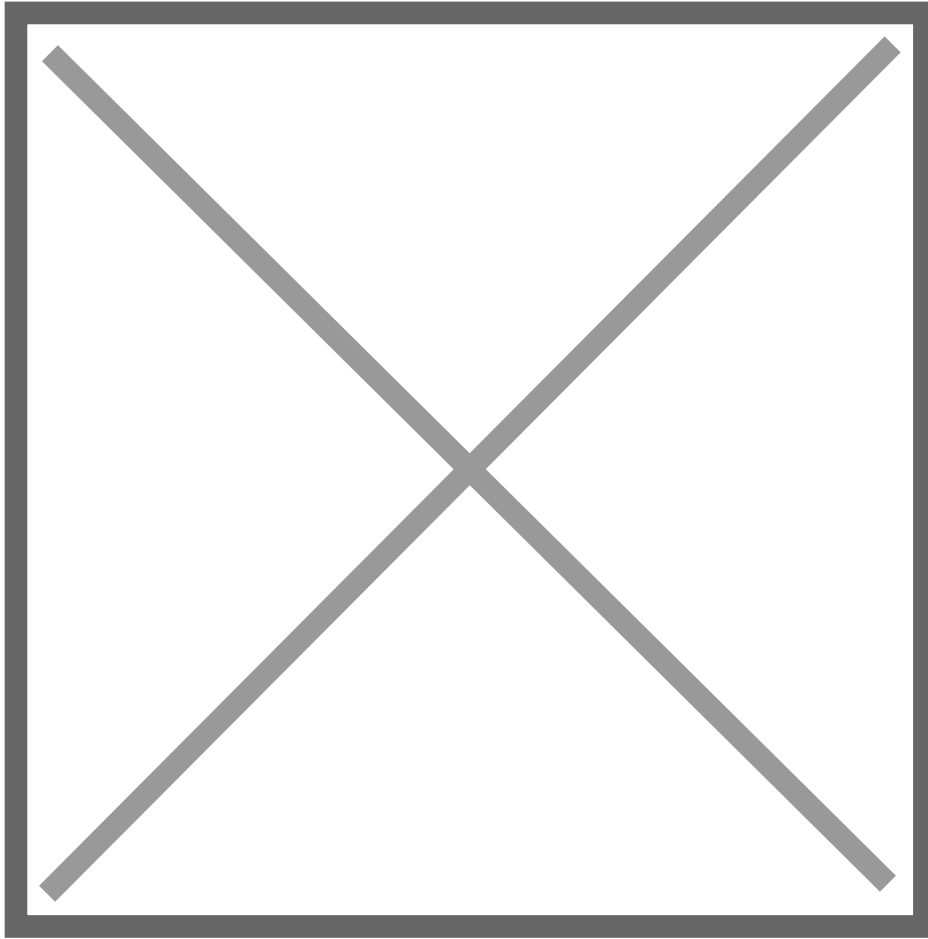
LLM.INT8()[] gemm[] X[] W
[] threshold[] 6
[] 6[]

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

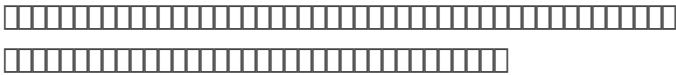




GPTQ ☐ OBS ☐



-



Group size

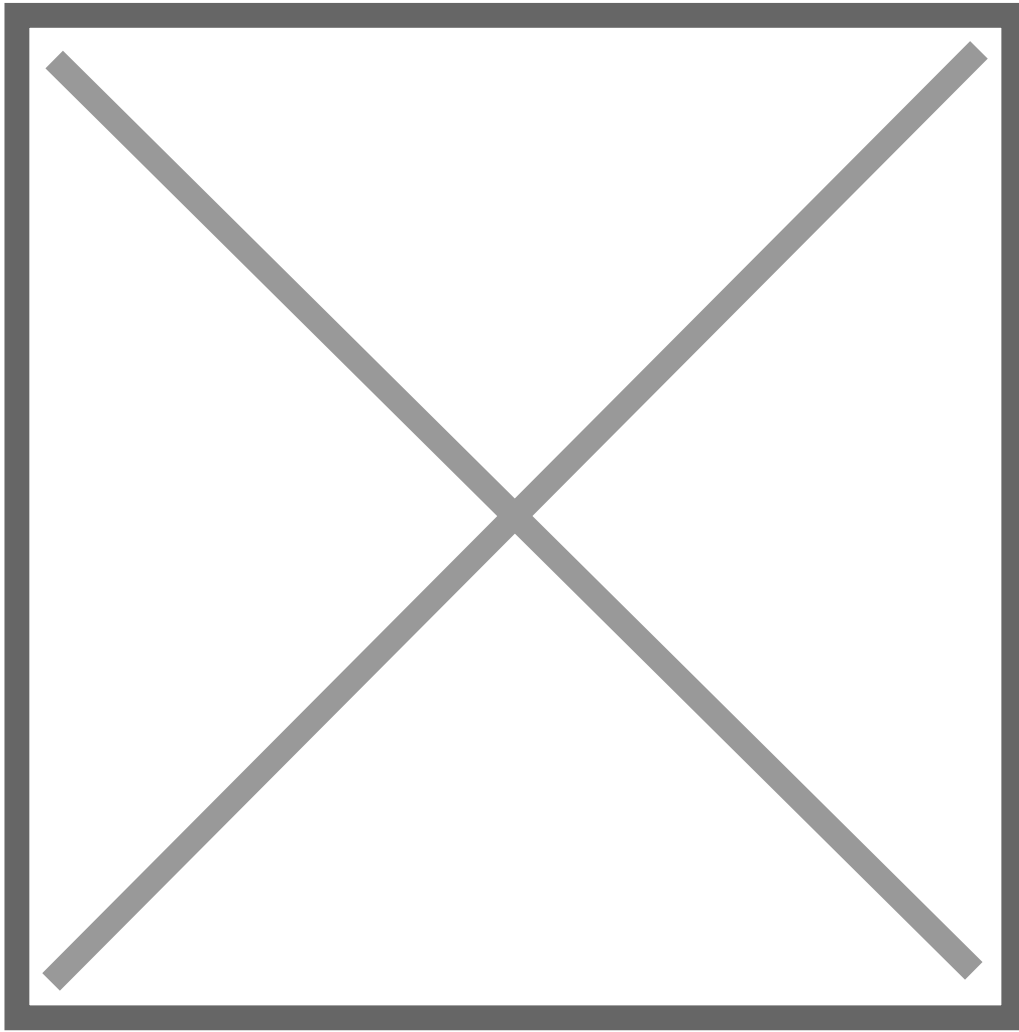


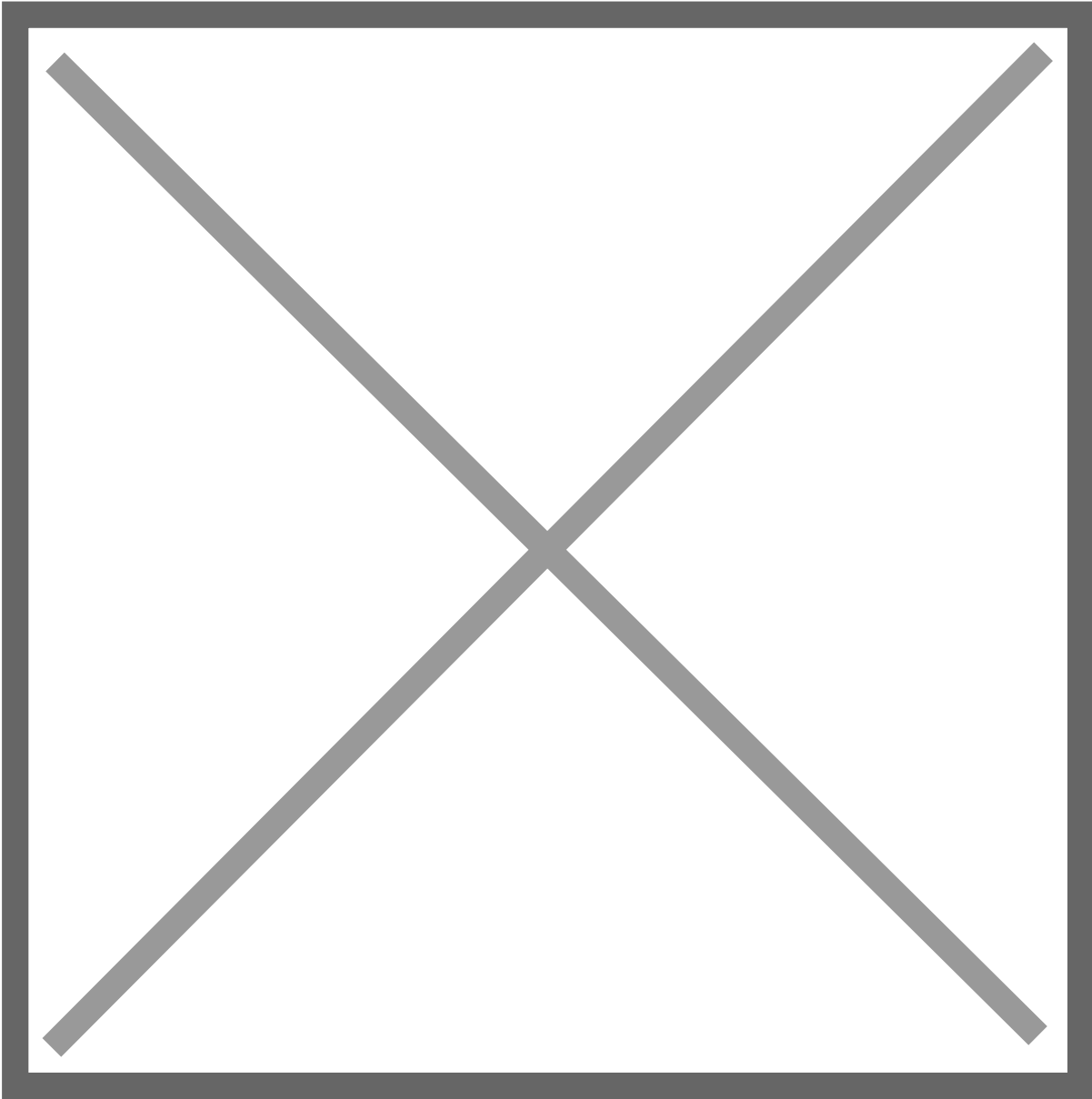
LLM

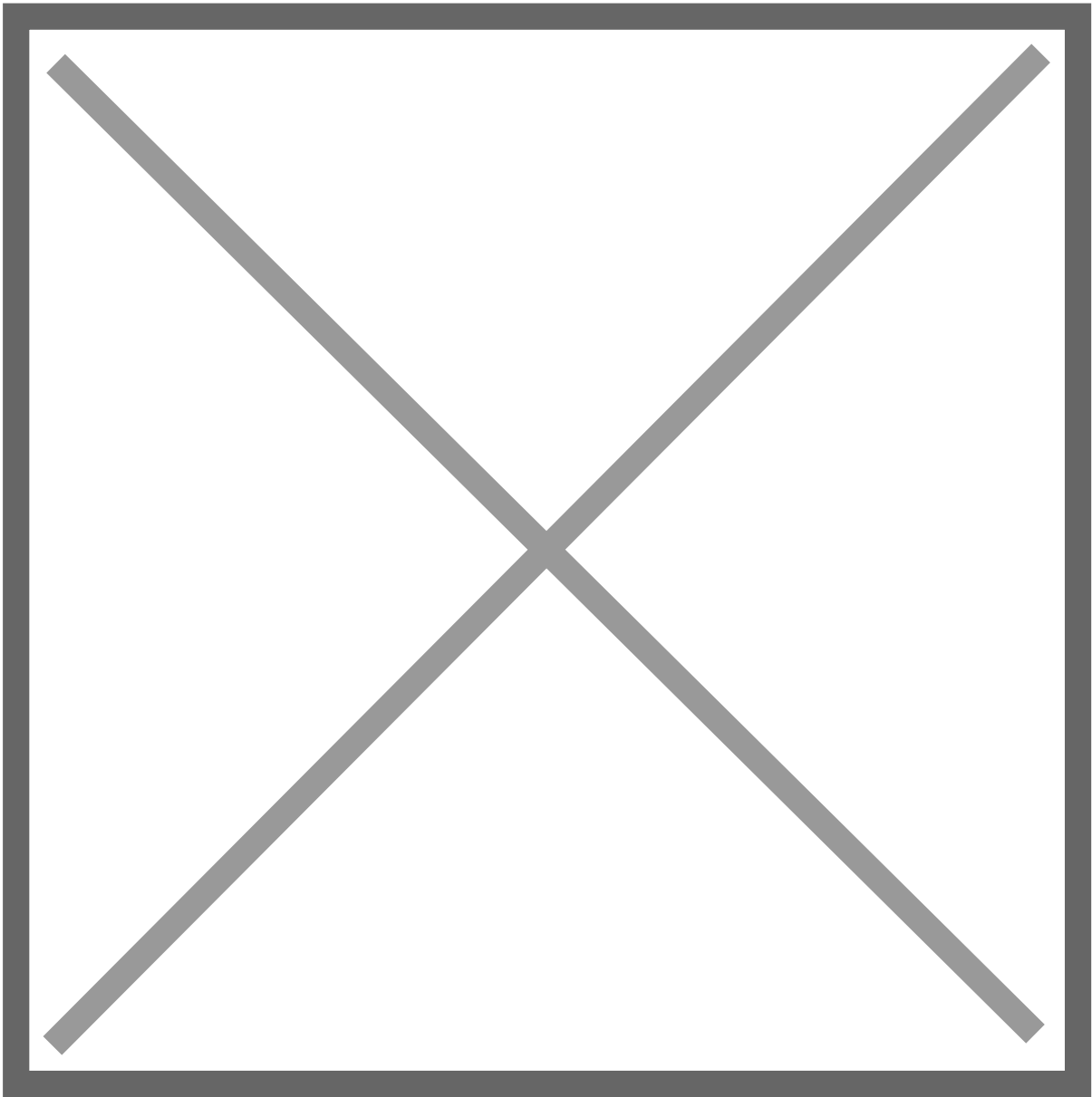


SmoothQuant: W8A8 scale

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

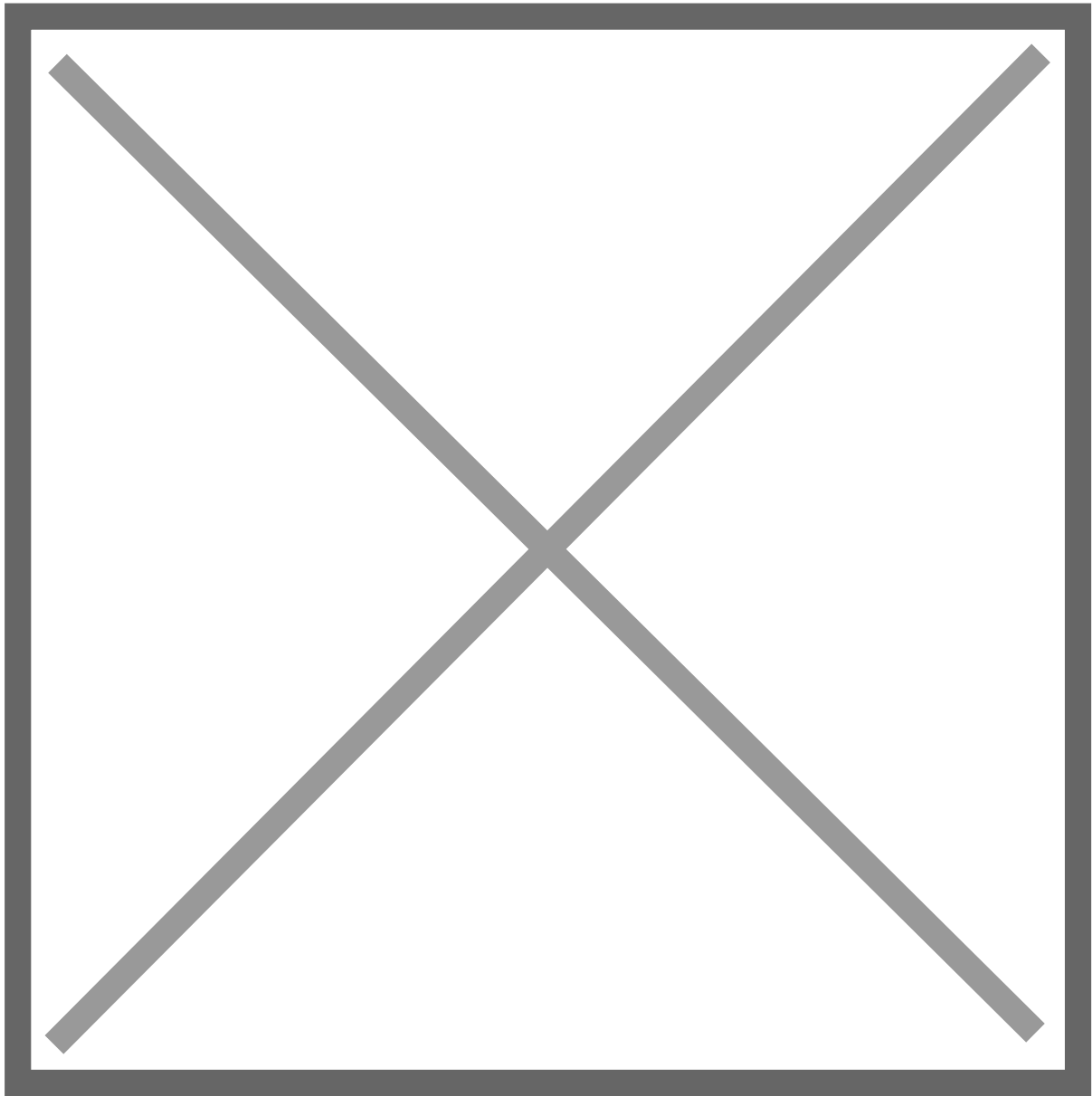




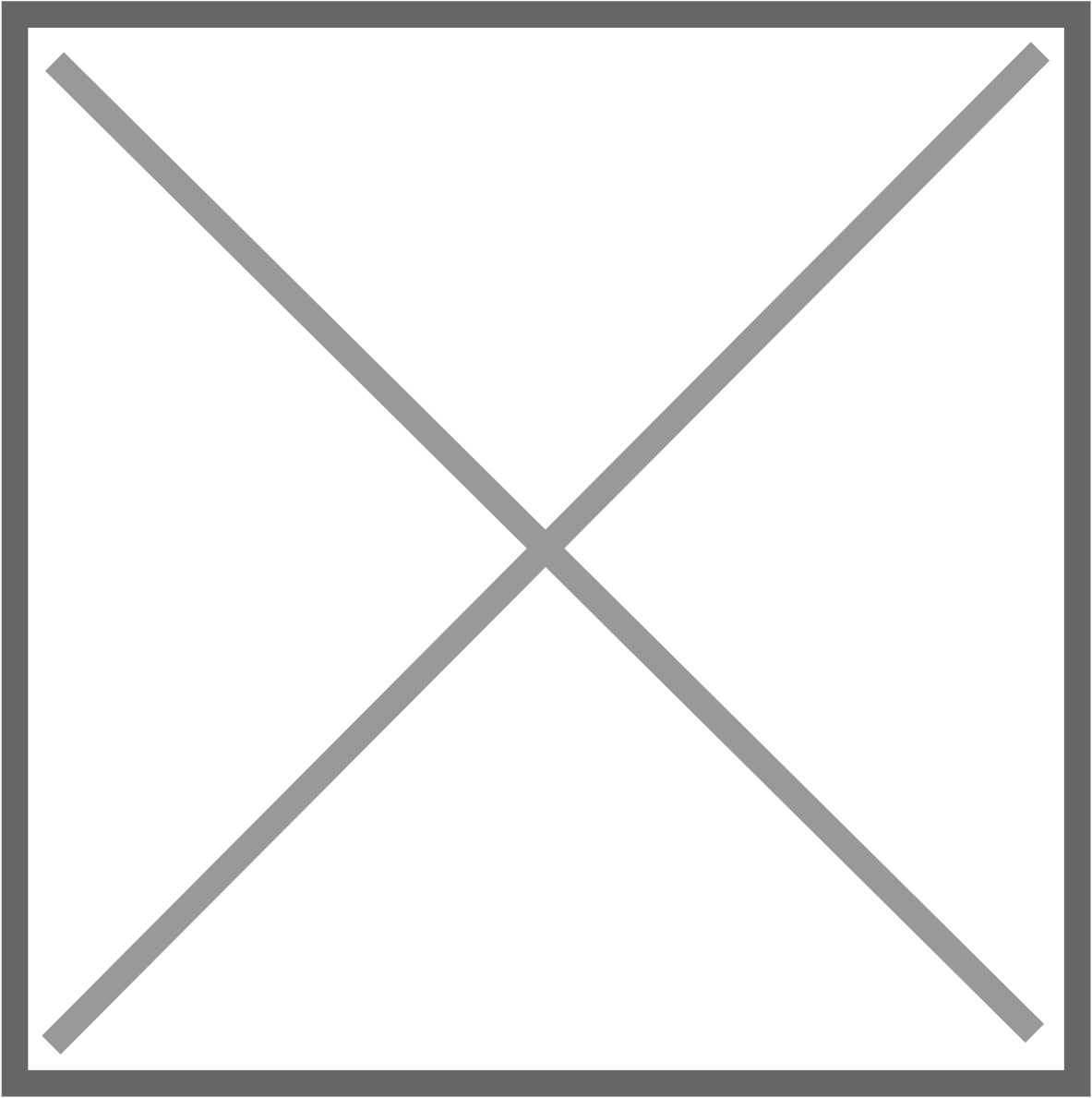


LLM weight-only SmoothQuant LLM
SmoothQuant W8A8
INT4 bit
SmoothQuant LLM
SmoothQuant LLM
weight-only
TensorRT-LLM TensorRT LLM
GPT2 SmoothQuant

AWQ:



The diagram illustrates the SmoothQuant process for quantizing LLM weights. It shows a 1D vector of 128 elements (8x16) being processed by SmoothQuant to produce a quantized tensor of 128 elements (8x16) with a 1% quantization error. The quantized tensor is then used to calculate the product of two 128x128 matrices (A and B) to produce a 128x128 result matrix (C).



□□□□□□□

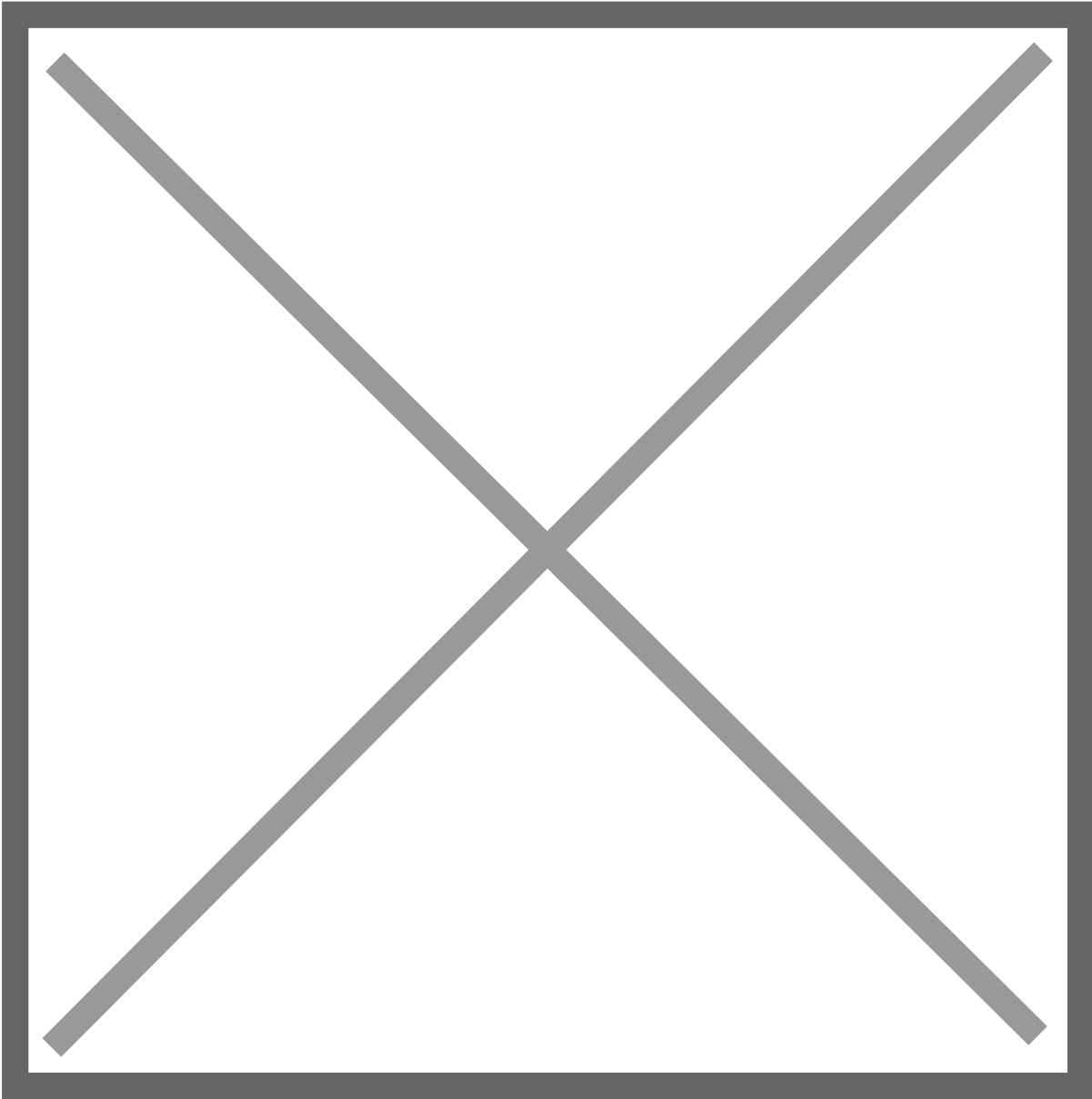
s□□□□□□□□

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

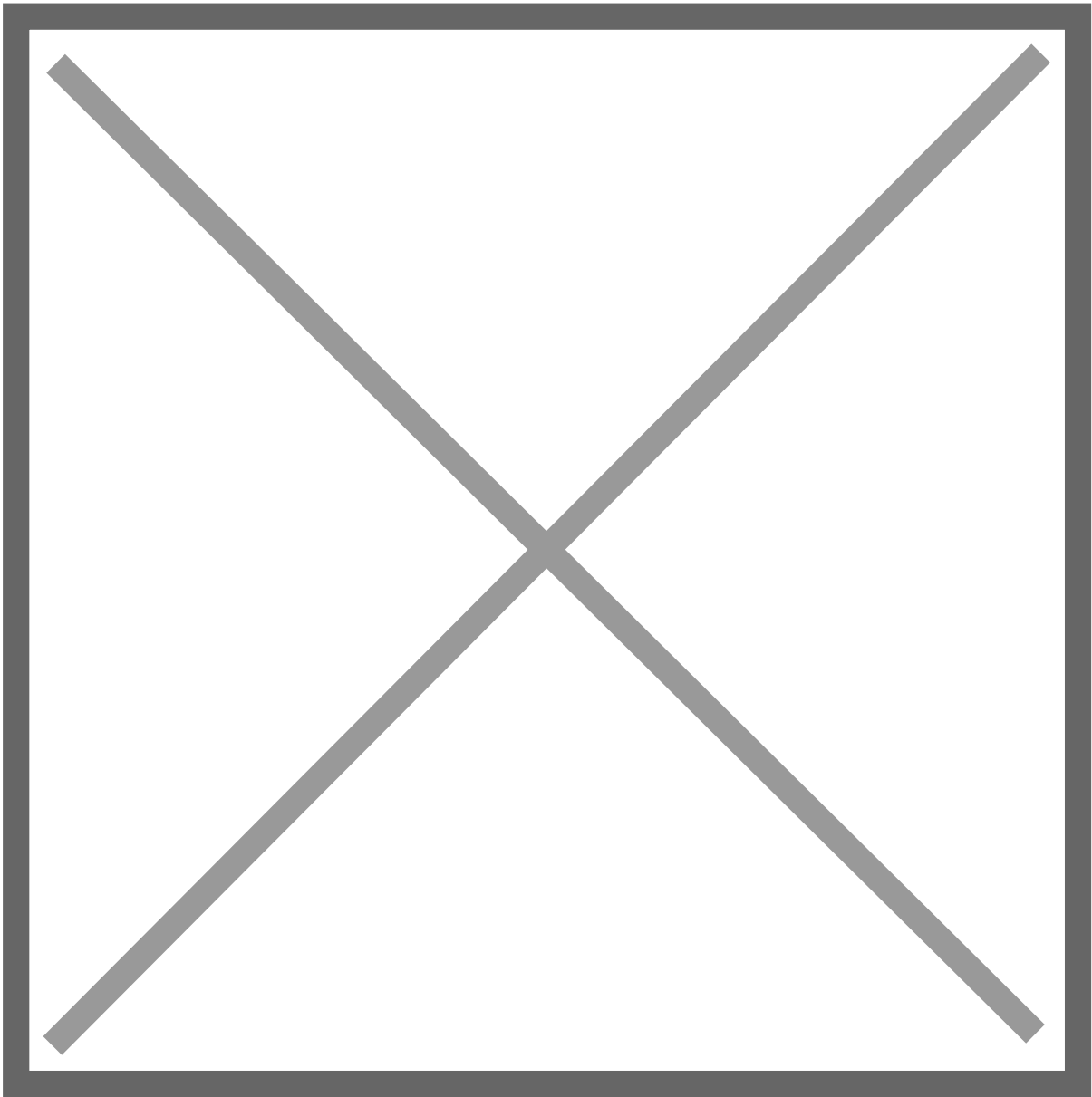
- Top Section:** Shows operations involving $sX = \text{meanc_out}[X]$, $sW = \text{meanc_out}[W]$, sX , group , weight , α , β , and $[0, 1]$. It also shows a sW vector and a group vector.
- Middle Section:** Shows a kernel vector and a weight vector.
- Bottom Section:** Shows operations involving smooth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



$y = K$ K-quant K-quant 16 x 8 “ ” 16
8 scale 16
1 fp16 K_2 16
“ ” size 16*8 “ ” K-
quant K_1 group size = 8, 65B 37.5GB

llama.cpp K-quant fp16 3~4 (
 token/s) fp16
 llama.cpp c++ cuda
 +gemm GPT
 KV Cache token vector-matrix GEVM
 matrix-matrix GEMM llama.cpp

llama.cpp llama +K-quant

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