

W4A8 Quantization: A Comparative Analysis for Efficient Neural Network Inference

Author: BenchChem Technical Support Team. **Date:** April 2026

Compound of Interest

Compound Name: *FPTQ*

Cat. No.: *B15621169*

[Get Quote](#)

A deep dive into 4-bit weight and 8-bit activation quantization methods, offering a comparative look at their performance, methodologies, and ideal use cases for researchers and scientists in drug development and other computationally intensive fields.

The relentless growth in the size and complexity of neural network models has spurred the development of various optimization techniques to reduce their computational and memory footprints. Among these, quantization, the process of reducing the precision of a model's weights and activations, has emerged as a powerful tool. This guide provides a comparative study of W4A8 quantization methods, where weights are quantized to 4 bits and activations to 8 bits, a scheme that strikes a compelling balance between model compression and accuracy.

Performance Benchmark

The following tables summarize the performance of different W4A8 quantization methods on widely used language models and benchmarks. These results highlight the trade-offs between various techniques and their effectiveness in preserving model performance post-quantization.

Perplexity Scores on Language Modeling Tasks

Perplexity is a standard metric for evaluating language models, with lower values indicating better performance. The table below compares the perplexity of LLaMA and OPT models quantized using different W4A8 schemes from the ZeroQuant-FP study. The baseline is the full-precision FP16 model.

Model	Dataset	FP16 (Baseline)	W8A8 (INT)	W4A8 (INT)	W8A8 (FP)	W4A8 (FP)
LLaMA-7B	Wikitext-2	5.03	5.21	5.61	5.14	5.48
PTB	29.4	30.2	32.5	29.8	31.7	
C4	7.61	7.82	8.31	7.75	8.15	
LLaMA-13B	Wikitext-2	4.45	4.61	4.95	4.55	4.86
PTB	25.9	26.6	28.5	26.3	27.9	
C4	6.92	7.11	7.53	7.04	7.41	
OPT-6.7B	Wikitext-2	11.9	12.5	13.8	12.2	13.2
PTB	58.1	61.2	66.7	59.9	64.1	
C4	14.3	15.1	16.5	14.7	15.8	
OPT-13B	Wikitext-2	11.2	11.8	12.9	11.5	12.4
PTB	54.2	57.1	61.9	55.8	59.7	
C4	13.5	14.2	15.4	13.8	14.8	

Data sourced from the ZeroQuant-FP research paper.

Accuracy on Common Sense Reasoning Tasks

The following table showcases the performance of the Fine-grained Post-Training Quantization (**FPTQ**) method on various common sense reasoning benchmarks for the LLaMA model family. Accuracy is reported as a percentage.

Model	Task	FP16 (Baseline)	SmoothQuant (W8A8)	Vanilla W4A8	FPTQ (W4A8)
LLaMA-7B	WinoGrande	73.74	73.78	64.18	73.80
PIQA	80.4	80.5	78.9	80.6	
HellaSwag	78.9	79.0	75.4	79.1	
ARCe	53.2	53.5	49.8	53.6	
LLaMA-13B	WinoGrande	76.19	76.36	69.90	75.74
PIQA	81.7	81.8	80.1	81.9	
HellaSwag	81.2	81.4	78.3	81.5	
ARCe	57.1	57.4	53.2	57.5	
LLaMA-65B	WinoGrande	79.20	78.69	69.98	79.14
PIQA	83.1	82.9	81.2	83.2	
HellaSwag	83.5	83.3	80.1	83.6	
ARCe	61.2	60.9	56.7	61.3	

Data sourced from the **FPTQ** research paper.[\[1\]](#)

Experimental Protocols

To ensure the reproducibility and clear understanding of the presented results, the following sections detail the experimental methodologies employed in the cited studies.

ZeroQuant-FP

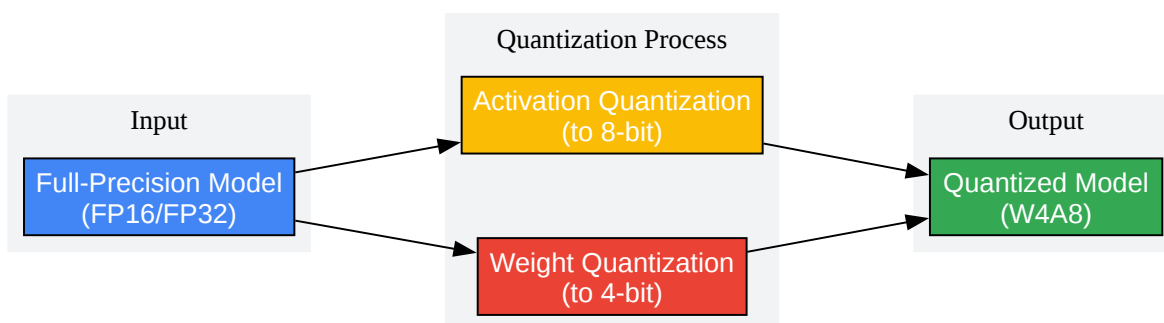
The ZeroQuant-FP experiments utilized a post-training quantization approach. For the calibration process, a random selection of 128 sentences from the C4 dataset was used. The performance evaluation was conducted on the Wikitext-2, PTB (Penn Treebank), and C4 datasets, with perplexity serving as the primary metric. The study compared both integer (INT) and floating-point (FP) representations for weights and activations within the W8A8 and W4A8 frameworks.

Fine-grained Post-Training Quantization (FPTQ)

The **FPTQ** method also follows a post-training quantization paradigm. The calibration set for **FPTQ** was constructed by randomly sampling 512 instances from the Pile dataset.^[1] The performance of the **FPTQ** W4A8 method was evaluated against a full-precision FP16 baseline, the SmoothQuant W8A8 method, and a naive W4A8 implementation. The evaluation was performed on the LAMBADA dataset for language modeling and a suite of common sense reasoning tasks, including WinoGrande, PIQA, HellaSwag, and ARCe, using the Language Model Evaluation Harness.^[1]

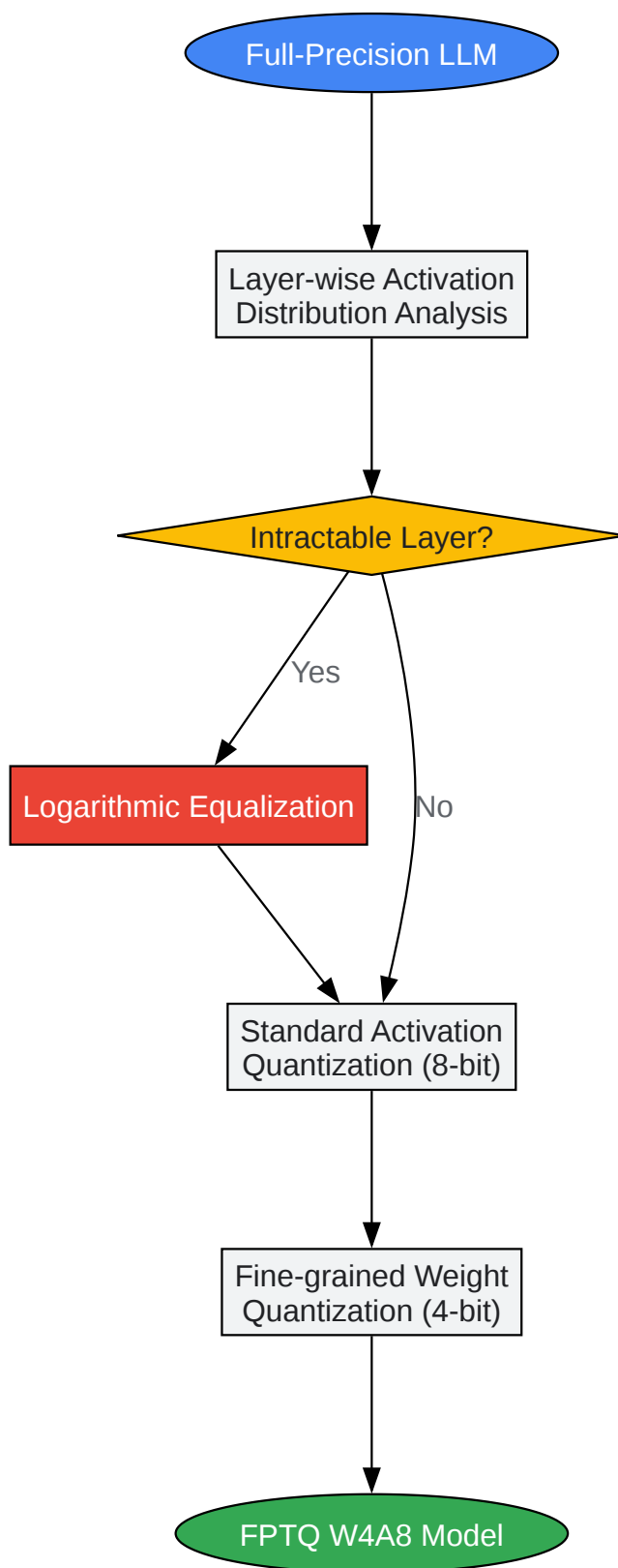
Visualizing Quantization Workflows

The following diagrams, generated using the DOT language, illustrate the conceptual workflows of W4A8 quantization.



[Click to download full resolution via product page](#)

A generalized workflow for W4A8 quantization.



[Click to download full resolution via product page](#)

Logical flow of the **FPTQ** method.

Need Custom Synthesis?

BenchChem offers custom synthesis for rare earth carbides and specific isotopic labeling.

Email: info@benchchem.com or [Request Quote Online](#).

References

- 1. openreview.net [openreview.net]
- To cite this document: BenchChem. [W4A8 Quantization: A Comparative Analysis for Efficient Neural Network Inference]. BenchChem, [2026]. [Online PDF]. Available at: [<https://www.benchchem.com/product/b15621169/docs#w4a8-quantization-a-comparative-analysis-for-efficient-neural-network-inference>]

Disclaimer & Data Validity:

The information provided in this document is for Research Use Only (RUO) and is strictly not intended for diagnostic or therapeutic procedures. While BenchChem strives to provide accurate protocols, we make no warranties, express or implied, regarding the fitness of this product for every specific experimental setup.

Technical Support: The protocols provided are for reference purposes. Unsure if this reagent suits your experiment?

Need Industrial/Bulk Grade? [Request Custom Synthesis Quote](#)

BenchChem

Our mission is to be the trusted global source of essential and advanced chemicals, empowering scientists and researchers to drive progress in science and industry.

Contact

Address: 3281 E Guasti Rd

Ontario, CA 91761, United States

Phone: (601) 213-4426

Email: info@benchchem.com

Contact our Ph.D. Support Team for a compatibility check