

QERA

An Analytical Framework For Quantization Error Reconstruction

Cheng Zhang
Jeffrey T. H. Wong, Can Xiao
A. Constantinides, Yiren Zhao



Quantization + Low-rank approximation

Compression

- Given a trained layer

$$y = xW$$

- Quantize

$$W_q = \text{quantize}(W)$$

- Then approximate error with low-rank terms:

$$(A_k, B_k) = \text{approx}(W - W_q)$$

Inference

$$y = xW_q + xA_kB_k,$$

W_q is **high-rank low-precision**,

A_k, B_k are **low-rank high-precision**.

Applications

- Quantized parameter-efficient fine-tuning (QPEFT)
- Post-training quantization (PTQ)

Existing Methods

Problem to solve

How to **approximate** $y = xW$ with $y = x(W_q + C_k)$?

- Ignore error: qLoRA
- Minimize weight error

$$\text{argmin}_{C_k} \|W - W_q - C_k\|_F^2$$

e.g, LoftQ, based on SVD on weights $C_k = \text{SVD}_k(W - W_q)$

- Minimize output error

$$\text{argmin}_{C_k} \mathbb{E}_{y \sim Y} \|y - \tilde{y}\|_2^2 \Rightarrow \text{argmin}_{C_k} \mathbb{E}_{x \sim X} \|x(W - W_q - C_k)\|_2^2$$

e.g, LQER, LQ-LoRA, based on heuristics

QERA

QERA provides closed-form solutions to minimizing output error

QERA-exact (exact solution)

$$C_k = \left(R_{XX}^{-\frac{1}{2}}\right)^{-1} \text{SVD}_k(R_{XX}^{-\frac{1}{2}}(W - W_q))$$

where $R_{XX} := \mathbb{E}_{x \sim X}[x^T x]$ is the auto-correlation matrix of input vector space

QERA-approx (approximate solution)

$$C_k = S^{-1} \text{SVD}_k(S(W - W_q))$$

where $S = \text{diag}(\sqrt{E[x \odot x]})$, $\odot$ denotes elementwise multiply.

Evaluation: Improved QPEFT & PTQ

QPEFT fine-tuning RoBERTa-base on GLUE

W-bits	Rank	Method	MNLI Acc	QNLI Acc	RTE Acc	SST Acc	MRPC Acc	CoLA Matt	QQP Acc	STS-B P/S Corr	Avg.
16	-	Full FT	87.61	92.95	73.16	94.88	92.15	60.41	91.61	90.44/90.25	85.38
16	8	LoRA	87.85	92.84	69.55	94.46	89.99	57.52	89.83	89.92/89.83	84
		QLoRA	87.21	92.32	63.9	94.08	88.24	56.08	90.55	89.59/89.56	82.75
4.25	8	LoftQ (5-iter)	87.27	92.48	67.13	94.38	88.24	54.59	90.51	88.75/88.79	82.92
		QERA-approx	87.28	92.45	70.4	94.38	88.97	55.99	90.39	89.83/89.72	83.71
		QLoRA	84.87	89.58	53.67	91.02	73.94	3.12	89.31	84.80/84.38	71.29
3.25	8	LoftQ (5-iter)	85.24	89.65	58.24	92.05	75.82	11	88.93	85.55/85.27	73.31
		QERA-approx	85.58	90.74	58.48	92.59	82.19	32.98	89.41	87.43/87.08	77.43
		QLoRA	77.87	85.26	54.15	90.02	71	0	87.93	74.72/75.31	67.62
2.5	64	LoftQ (5-iter)	80.15	87.65	52.95	90.94	74.35	3.43	89.17	82.76/82.90	70.18
		QERA-exact	84.64	90.05	58.48	92.32	84.72	26.43	89.69	86.48/86.40	76.23

PTQ perplexity (↓) across various models

W-bits	Rank	Method	TinyLlama 1.1B	Gemma-2 2B	Phi-3.5 3.8B	LLaMA-2 7B	LLaMA-2 13B	LLaMA-3.1 8B	LLaMA-3.1 80B
-	-	BF16	13.98	13.08	11.5	8.71	7.68	7.55	3.06
4.25	-	HQQ	15.02	14.29	14.63	9.59	8.27	8.72	3.97
		w-only	19.4	16.23	14.16	9.45	8.06	8.78	4.55
		ZeroQuant-V2	18.03	15.71	14.09	9.42	8.07	8.83	4.48
4.25	32	LQER	16.23	14.55	1				