

RQT: Hierarchical Residual Quantization for Multi-Model Compression

Anonymous ACL submission

Abstract

Delta compression methods focus on efficiently serving multiple uniquely fine-tuned models, each tailored to specific tasks and user requirements. These approaches decompose a fine-tuned LLM into a base model and corresponding delta weights, which are compressed using low-rank or low-bit representations to reduce storage costs. However, their effectiveness is highly sensitive to the magnitude of the model deltas—a factor directly influenced by the scale of the training data. We propose the **Residual Quantization Tree (RQT)**, a hierarchical quantization framework that automatically shares low-bit integer weights across similar fine-tuned models. The RQT construction employs a two-phase greedy algorithm: a bottom-up aggregation of models based on weight matrix similarity and top-down residual quantization, in which each node optimizes the quantization parameters and then delegates residual errors to child nodes. We evaluate RQT on fine-tuned models across mathematics, coding, chatbot, and Chinese LLMs. The results show that RQT achieves an average accuracy degradation of approximately 3% (comparable to previous 4-bit post-training quantization) while maintaining an effective bitwidth of around 2 bits.

1 Introduction

Large Language Models (LLMs), such as GPT (OpenAI, 2023b) and LLaMA (Touvron et al., 2023), have demonstrated exceptional performance and are widely applied in various applications, including chatbots (Chiang et al., 2023; OpenAI, 2023b; Touvron et al., 2023) and coding assistants (Luo et al., 2024; Wei et al., 2024; Rozière et al., 2023). These models achieve high accuracy in target domains through a two-stage process: pre-training on a large corpus of text data, followed by fine-tuning on task-specific datasets, such as code (Luo et al., 2024), math (Xiong et al., 2024), or

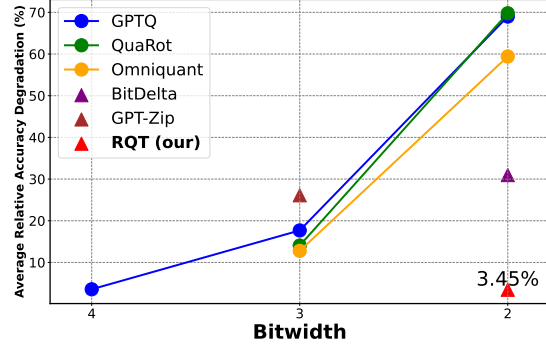


Figure 1: We evaluate RQT on fine-tuned models across mathematics, coding, chatbot, and Chinese LLMs. The results show that RQT achieves an average accuracy degradation of approximately 3% (comparable to previous 4-bit post-training quantization) while maintaining an effective bitwidth of around 2 bits.

human preferences (OpenAI, 2023b). Cloud AI infrastructure providers, such as OpenAI (OpenAI, 2023a), Google (Google, 2023), and Microsoft (Microsoft, 2023) offer APIs that allow users to fine-tune pre-trained LLMs with their own data, enabling the creation of customized model variants. This paradigm has driven substantial research into developing efficient methods for serving millions of uniquely fine-tuned models, each tailored to individual tasks and user requirements (Liu et al., 2024a; Ping et al., 2024; Yao and Klimovic, 2023; Ryu et al., 2023).

Although systems like Punica (Chen et al., 2024) and S-LoRA (Sheng et al., 2023) can scale to serve thousands of Low-Rank Adaptation (LoRA) adapters, these methods are incompatible with traditional full-parameter fine-tuning approaches. Notably, most models on the Open LLM Leaderboard (Fourrier et al., 2024) still rely on full-parameter fine-tuning due to the performance gap between full-parameter fine-tuning and LoRA (Biderman et al., 2024; Ding et al., 2023). To address these challenges, delta compression methods (Liu et al.,

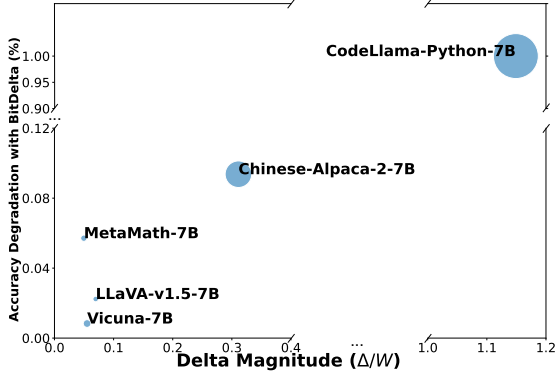


Figure 2: This figure illustrates the relationship between relative accuracy degradation and delta magnitude across models for various tasks. The circle size represents the training dataset size.

2024a; Ping et al., 2024; Yao and Klimovic, 2023; Ryu et al., 2023) decompose the fine-tuned model weights W_{fine} into the weights of the base pre-trained model W_{base} and the delta induced by the fine-tuning process Δ ($\Delta = W_{fine} - W_{base}$). This delta is then compressed using techniques such as quantization (Liu et al., 2024a; Isik et al., 2023), pruning (Yu et al., 2024), low-rank matrix factorization (Ping et al., 2024; Ryu et al., 2023), or combinations of these methods.

Delta compression methods are based on the observation that model deltas generally have smaller magnitudes and inherent sparsity (Yu et al., 2024; Liu et al., 2024a; Yao and Klimovic, 2023). As a result, deltas tend to be more compressible than the original model weights. For example, BitDelta (Liu et al., 2024a) achieves 1-bit quantization of delta weights, while DARE (Yu et al., 2024) can prune up to 90% of delta weights. However, this assumption primarily applies when fine-tuning on small datasets, such as instruct tuning (Chiang et al., 2023; Touvron et al., 2023; Xiong et al., 2024; Liu et al., 2023a). As shown in Figure 2, as the size of the training data increases, the relative magnitude of the delta $|\Delta|/|W_{fine}|$ also increases, leading to a greater degradation of accuracy when applying delta compression methods to Chinese and Code LLM. Consequently, BitDelta confines its experiments to chat models, such as LLaMA-2-chat (Touvron et al., 2023), Vicuna (Chiang et al., 2023), and WizardLM (Xu et al., 2024). In contrast, while Delta-CoMe (Ping et al., 2024) broadens its scope by evaluating delta compression methods across various tasks, including mathematics, coding, chat, and multi-modal applications, it still faces a lim-

itation: it selects different base models for each task to ensure that the fine-tuned models remain sufficiently close to their base models.

This paper proposes a novel quantization framework, the **Residual Quantization Tree (RQT)**, for efficiently serving multiple fine-tuned models. An RQT is a hierarchical structure in which the nodes share low-bit integer weights between similar models. As illustrated in Figure 3, models are recursively partitioned into child nodes until reaching the leaf nodes, each corresponding to an individual model. The quantized weights are computed by aggregating the quantized values at each node along the path from the root to the corresponding leaf. To construct an RQT, we introduce a practical two-step greedy approach. (a) Bottom-up tree construction: Starting from leaf nodes (individual models), we iteratively merge node pairs with minimal weight matrix divergence (measured by L_2 distance), forming parent nodes until reaching the root. (b) Top-down residual quantization: Beginning at the root, we compute optimal low-bit integer weights for contained models at each node, then propagate residual errors to child nodes for refinement, recursively minimizing the weight reconstruction error. Our contributions are threefold:

- We propose the **Residual Quantization Tree (RQT)**, which addresses the sensitivity of previous delta compression methods to the magnitude of delta values.
- To construct an RQT, We propose an efficient two-step greedy approach that decouples model assignment and residual quantization.
- Experiments demonstrate the effectiveness of RQT in comparison with previous PTQ and delta-compression methods. Even with an average bitwidth of approximately 2 bits, RQT achieves an average accuracy degradation of approximately 3% on fine-tuned LLaMA-2-7B models across various tasks.

2 Related Works

2.1 Model Quantization

Quantization has been extensively applied to accelerate model inference (Jacob et al., 2018; Nagel et al., 2020; Li et al., 2021), and it has become a crucial technique for LLMs in the current era of rapid development (Wei et al., 2022; Xiao et al., 2023; Lin et al., 2023; Frantar et al., 2022; Liu

et al., 2023b). Depending on whether activations are quantized, quantization can be categorized into weight-only quantization and weight-activation quantization. Weight-only quantization reduces memory usage and inference latency by quantizing weights into low-bit precision while keeping activations in floating-point. For example, GPTQ (Frantar et al., 2022) introduces a layer-wise quantization approach based on approximate second-order information. OmniQuant (Shao et al., 2023) incorporates learnable parameters to clip extreme weight values and shift the quantization challenge from activations to weights. Recently, QuaRot (Ashkboos et al., 2024) and SpinQuant (Liu et al., 2024b) facilitate quantization by rotating LLMs to eliminate outliers in activations without affecting the output. Despite these advancements, ultra-low-bit quantization (e.g., 2-bit) continues to experience significant performance degradation.

2.2 Delta Compression

The rise of millions of uniquely fine-tuned models has recently driven significant research interest in delta compression methods. Specifically, delta compression involves subtracting the base pre-trained model weights from fine-tuned weights to obtain the model delta. Since model deltas generally exhibit smaller magnitudes and inherent sparsity, they are often more compressible. For example, GPT-Zip (Isik et al., 2023) quantizes model deltas into 2/3/4 bits using the post-training quantization (PTQ) method GPTQ (Frantar et al., 2022). DARE (Yu et al., 2024) eliminates the majority of delta weights (up to 90%) with minimal impact on the performance of aligned LLMs. BitDelta (Liu et al., 2024a) binarizes model deltas and fine-tunes the quantization scales through knowledge distillation. Delta-CoMe (Ping et al., 2024) applies singular value decomposition (SVD) to deltas and assigns varying bitwidths to different singular vectors based on their singular values.

Although these methods achieve high compression ratios by quantizing deltas, they face limitations when applied to models across diverse tasks. In other words, the assumption that model deltas are smaller in magnitude than the original weights is not always valid, especially when fine-tuning with large datasets, such as those for code and Chinese language models.

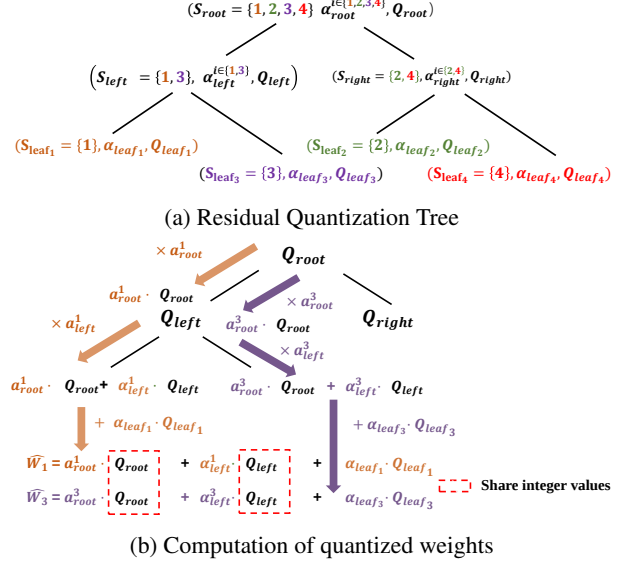


Figure 3: An illustration of an RQT with four fine-tuned models. Different colors denote each model. The quantized weights $\hat{W}_1$ and $\hat{W}_3$ are computed by summing along the path from the root to the leaf. They use the shared integer values Q_{root} and Q_{left} in their common ancestor node.

3 Methods

3.1 Formulation

Let $\{W_i \in \mathbb{R}^{d_o \times d_i}\}_{i=1}^N$ represent the weight matrices of a target linear layer across N fine-tuned models, where d_i and d_o represent the input and output dimensions, respectively. We define the Residual Quantization Tree (RQT) as follows:

Definition 1 (Residual Quantization Tree). Given a set of weight matrices $\{W_i \in \mathbb{R}^{d_o \times d_i}\}_{i=1}^N$ from N fine-tuned models, a Residual Quantization Tree is a tree where each node v contains:

- $S_v \subseteq \{1, \dots, N\}$: The index set of assigned models to node v .
- $Q_v \in \mathbb{Z}^{d_o \times d_i}$: A *shared* integer matrix for quantizing *all* weight matrices assigned to node v , i.e., $\{W_i \mid i \in S_v\}$;
- $\Sigma_v = \{\alpha_v^i \in \mathbb{R}^{d_o} \mid i \in S_v\}$: The *model-specific* quantization scales, where each α_v^i is specific to W_i at node v ;
- $C_v = \{v_1, v_2, \dots, v_m\}$: The child nodes of v satisfying $\bigcup_{j=1}^m S_{v_j} = S_v$ with $S_{v_j} \cap S_{v_k} = \emptyset$

The root node encompasses all models with $S_{\text{root}} = \{1, \dots, N\}$. Every non-leaf node v distributes its assigned models S_v to child nodes C_v until the leaf nodes are reached, each containing a

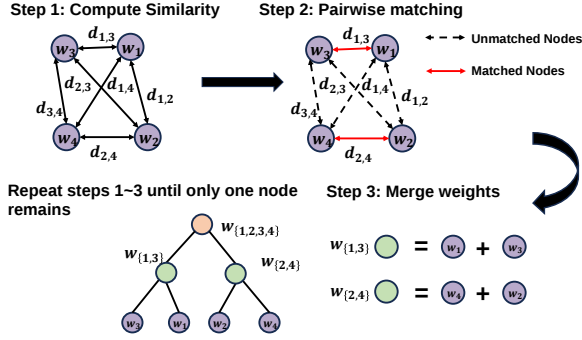


Figure 4: Illustration of the bottom-up construction process of the RQT.

single model ($|S_{\text{leaf}}| = 1$). At each node v , weight matrix W_i is quantized as $\alpha_v^i \cdot Q_v$, using the shared low-bit values Q_v and the individual per-channel quantization scales α_v^i . This architecture achieves higher effective bitwidth than the average storage bitwidth, establishing superior compression-performance trade-offs compared to conventional quantization methods.

Below, we describe how to compute the quantized weights given the RQT:

Definition 2. For weight matrix W_i , its quantized weights $\widehat{W}_i$ are computed through summation along the path from the root to the corresponding leaf:

$$\widehat{W}_i = \sum_{v \in \mathcal{P}(i)} \alpha_v^i \cdot Q_v \quad (1)$$

where $\mathcal{P}(i)$ denotes the path from root to leaf node contains i -th model ($S_{\text{leaf}} = \{i\}$).

Remarks. RQT fundamentally advances delta compression through its hierarchical residual decomposition. Unlike conventional approaches that decompose model weights as $\Delta = W_{\text{fine}} - W_{\text{base}}$ with manually selected W_{base} , the RQT has two advantages: (a) The hierarchical representation of inter-model similarities aligns naturally with multi-stage LLM fine-tuning pipelines. (b) RQT achieves adaptive bitwidth allocation compared to delta compression, making it less sensitive to delta magnitude and more robust to task diversity.

3.2 Algorithm

In this section, we propose a greedy two-step approach for constructing an RQT. For simplicity, we assume that the tree is binary ($|C_v| = 2$ for all nodes) and that the number of models N is a power of 2. In terms of bit allocation, we apply binary quantization (1-bit) to all nodes except the

root node, which is allocated higher bits (e.g., 4-bits). This simplification yields a complete binary tree with an average bitwidth of $\frac{4+(2N-2) \times 1}{N} \approx 2$ bits. Section 4.4 later relaxes these assumptions to handle arbitrary N by incorporating new models into an existing RQT.

3.2.1 Step 1: Tree Construction

Intuitively, models with higher similarity should be positioned closer in the RQT to enable them to have more common ancestor nodes. For example, in Figure 3, the weight matrices W_1 and W_3 should originate from similar models (e.g., both are math LLMs), as they utilize the same low-bit matrix Q_{left} in the second layer. The same applies to W_2 and W_4 .

To achieve this, we propose a greedy bottom-up tree construction method based on the similarity of weight matrices. An illustration of this procedure is provided in Figure 4. Specifically, we start with N leaf nodes $\{v_i \mid i \in 1, \dots, N\}$, each containing a single model $S_{v_i} = \{i\}$. Next, we pair nodes by minimizing the total distance between their respective weight matrices:

$$\begin{aligned} \min_{\{x_{ij}\}_{1 \leq i < j \leq N}} \sum_{1 \leq i < j \leq N} x_{ij} \cdot d_{ij}, \\ \text{subject to: } \sum_{j \neq i} x_{ij} = 1, \quad x_{ij} \in \{0, 1\}. \end{aligned} \quad (2)$$

Here, x_{ij} is a binary variable indicating whether nodes i and j are paired, and d_{ij} represents the distance between the weight matrices W_i and W_j , chosen as the L_2 distance in this work:

$$d_{ij} = \|W_i - W_j\|_2. \quad (3)$$

This problem can be solved using dynamic programming or linear integer programming techniques, as detailed in (Cormen et al., 2022). Once the optimal pairs $\{x_{ij}\}_{1 \leq i < j \leq n}$ are determined, we merge the matched nodes into their parent nodes by averaging their weights:

$$\begin{aligned} S_{\text{parent}(v_i, v_j)} &= S_{v_i} \cup S_{v_j}, \\ C_{\text{parent}(v_i, v_j)} &= \{v_i, v_j\}, \\ W_{\text{parent}(v_i, v_j)} &= \frac{1}{2}(W_i + W_j). \end{aligned} \quad (4)$$

This iterative process—which comprises distance computation, matching, and merging—continues until a single root node remains. The complete pseudocode is formally described

Algorithm 1 Tree_Construction

```

1: Input: A set of  $N$  weight matrices  $\Theta = \{W_i \mid i = 1, \dots, N\}$ 
2: Output: An incomplete residual quantization tree  $T$  with defined  $S$  and  $C$ 
3: Let  $\Phi = \{v_i \mid i = 1, \dots, N\}$  be the set of leaf nodes.
4: for  $i = 1$  to  $N$  do
5:   Initialize  $S[v_i] = \{i\}$ ,  $C[v_i] = \emptyset$ 
6: end for
7: while  $|\Phi| > 1$  do
8:   Compute  $x_{ij}$  according to Equation 2
9:   for each pair of nodes  $(v_i, v_j) \in \Phi \times \Phi$  do
10:    if  $x_{ij} = 1$  then
11:      Create a new parent node  $v_{\text{new}}$ 
12:       $W_{\text{new}} \leftarrow \frac{1}{2}(W_i + W_j)$ 
13:       $S[v_{\text{new}}] \leftarrow S[v_i] \cup S[v_j]$ 
14:       $C[v_{\text{new}}] \leftarrow \{v_i, v_j\}$ 
15:      Add  $v_{\text{new}}$  to  $\Phi$  and  $W_{\text{new}}$  to  $\Theta$ 
16:      Remove  $v_i$  and  $v_j$  from  $\Phi$ 
17:      Remove  $W_i$  and  $W_j$  from  $\Theta$ 
18:    end if
19:   end for
20: end while
21: return the root node in  $\Phi$ 

```

in Algorithm 1. Crucially, the newly generated merged weights are exclusively employed for similarity computation during tree construction. Upon completing the initial phase, we obtain an incomplete RQT with established model indices S_v and child nodes mapping C_v . The subsequent phase involves computing quantization parameters Q_v and Σ_v for each node.

3.2.2 Step 2: Residual Quantization.

In the second step, we apply residual quantization in a top-down manner. Specifically, we compute the optimal shared binary weights Q and quantization scales $\Sigma = \{\alpha_i \mid i \in S\}$ by minimizing the MSE for each node. Without loss of generality we assume W and Q are vectors in $\mathbb{R}^n$, where $n = d_i$ in per-channel quantization:

$$\begin{aligned}
J(Q, \Sigma) &= \sum_{i \in S} \|W_i - Q \cdot \alpha_i\|_2^2 \\
&= \sum_i \alpha_i^2 Q^T Q - 2\alpha_i Q^T W_i + W_i^T W_i
\end{aligned} \tag{5}$$

This objective differs from those of previous binary neural networks (Rastegari et al., 2016; Courbariaux and Bengio, 2016) in that the binary weights are shared across multiple models. To address this, we propose an iterative method to compute the optimal values for α_i^* and Q^* using Equation 6 and Equation 7, respectively. Specifically, the optimal solution for α_i is obtained by setting the partial derivative of the objective function with

Algorithm 2 Residual_Quantization

```

1: Input: A set of weights  $\Theta = \{W_i \mid i = 1, \dots, N\}$ 
2: Input: An incomplete residual quantization tree  $T$  with only defined  $C$  and  $S$ 
3: Output: A complete residual quantization tree  $T$  with defined  $Q$  and  $\Sigma$ 
4: while not converged do
5:   Compute  $\alpha_i$  using Equation 6
6:   Compute  $Q$  using Equation 7
7: end while
8: Update  $W_i \leftarrow W_i - Q \cdot \alpha_i$  for all  $i$ 
9: for each  $T_j \in C[T]$  do
10:   Residual_Quantization( $\{W_i \mid i \in S[T_j]\}$ ,  $T_j$ )
11: end for
12: return  $T$ 

```

respect to α_i to zero.:

$$\alpha_i^* = \frac{Q^T W_i}{Q^T Q} \tag{6}$$

Similarly, when fixed scales α_i , the optimal solution for Q can be derived by taking the sign of the weighted summation of all weights:

$$Q^* = \text{sign} \left(\sum_i \alpha_i W_i \right) \tag{7}$$

Equations 6 and 7 are applied alternately until convergence is reached. Next, the residual weights are computed and passed to the child nodes. The complete procedure is formally described in Algorithm 2.

$$W_i \leftarrow W_i - \alpha_i \cdot Q \tag{8}$$

4 Experiments

4.1 Experiment Setup

Tasks and Models. We conduct comprehensive evaluations across **four representative application domains** of aligned LLMs: **Mathematical Reasoning**, **Code Generation**, **Chatbots**, and **Chinese Language Understanding**. For each domain, we systematically select **four open-source models** fine-tuned on the LLaMA-2-7B model, creating a 16-model benchmark to establish a five-level RQT.

- **Mathematical Reasoning:** Employing GSM8K (Cobbe et al., 2021) and MATH (Hendrycks et al., 2021) benchmarks, we assess MetaMath-7B-V1.0 (Xiong et al., 2024), WizardMath-V1.0 (Luo et al., 2023), Xwin-Math-7B-V1.1 (Li et al., 2024), and MuggleMath-7B (Li et al., 2023a).

- **Code Generation:** Assessing Pass@1 (%) performance on HumanEval (Chen et al., 2021) and MBPP (Austin et al., 2021), we assess CodeLLaMA-7B-Python (Rozière et al., 2023), WizardCoder-Python-7B-V1.0 (Luo et al., 2024), Magicoder-S-CL-7B (Wei et al., 2024), and ReflectionCoder-CL-7B (Ren et al., 2024).
- **Chatbots:** Measuring accuracy on ARC-Easy (Clark et al., 2018) and HellaSwag (Zellers et al., 2019), we assess LLaMA-2-7B (Touvron et al., 2023), LLaMA-2-chat-7B (Touvron et al., 2023), Vicuna-v1.5 (Chiang et al., 2023), and Vicuna-7B-v1.5-16k (Chiang et al., 2023).
- **Chinese Language Processing:** Utilizing C-Eval (Huang et al., 2023) and CMMLU (Li et al., 2023b), we assess Chinese-Alpaca-2-7B (Cui et al., 2023), Chinese-LLaMA-2-7B (Cui et al., 2023), and their 16k variants.

Setup. We employ per-channel binary quantization (1-bit) for non-root nodes with 4-bit precision reserved for the root node, achieving an effective bitwidth of $\frac{4+30}{16} = 2.125$. Building upon BitDelta (Liu et al., 2024a), we optimize quantization scales α_i through end-to-end training while keeping the low-bit model weights frozen. This approach effectively incorporates activation distribution characteristics while significantly reducing GPU memory consumption through weight freezing. The calibration process utilizes 800 randomly sampled 128-token sequences from the C4 corpus (Pal et al., 2023). We perform 3 training epochs using the AdamW optimizer with a learning rate of 5×10^{-6} , executed on a single NVIDIA A800 GPU (80GB). The complete distillation process requires approximately 20 minutes for the 7B parameter model. Furthermore, inspired by QuaRot (Ashkboos et al., 2024) and SpinQuant (Liu et al., 2024b), we implement pre-quantization orthogonal rotation on weight matrices to enhance quantization robustness while introducing no additional inference overhead.

Baselines. We evaluate our framework against two categories of compression approaches: (1) PTQ methods including GPTQ (Frantar et al., 2022), OmniQuant (Shao et al., 2023), and QuaRot (Ashkboos et al., 2024); and (2) delta compression techniques represented by GPT-Zip (Isik et al., 2023) and BitDelta (Liu et al., 2024a). For

PTQ baselines, we employ group-wise quantization with a group size of 128, achieving an average storage overhead of 0.125 bits per element. For delta compression methods, we implement BitDelta and GPT-Zip using the LLaMA-2-7B (16-bit floating-point) as the base model, quantizing 16 model deltas to 1-bit and 2-bit precision, respectively. This configuration yields effectively 2 bits (BitDelta: $\frac{16 \times 1 + 16}{16}$) and 3 bits (GPT-Zip: $\frac{16 \times 2 + 16}{16}$) per element when accounting for base model storage.

4.2 Main Results

Comparison with PTQ Methods. Table 1 presents aggregated results across two evaluation datasets for LLaMA2-7B fine-tuned variants (more detailed results are provided in the Appendix). Our framework employs per-channel binary quantization (1-bit) for non-root nodes with 4-bit precision reserved for the root node, achieving an effective bitwidth of $\frac{4+30}{16} = 2.125$ bits/parameter – equivalent to conventional 2-bit group-wise quantization with group size 128 in storage cost. Remarkably, despite this comparable bandwidth, our method surpasses existing 3-bit PTQ approaches in accuracy, particularly on complex multi-step reasoning tasks (mathematical problem-solving and code generation). For instance, in mathematical reasoning evaluation, standard PTQ methods exhibit catastrophic degradation with accuracy dropping to nearly zero, while our approach maintains robust performance with a maximum degradation of only 2.51%.

Comparison with Delta Compression Methods. Figure 5 and Table 1 demonstrate the superiority of RQT over delta compression baselines (BitDelta (Liu et al., 2024a), GPT-Zip (Isik et al., 2023)) when models originate from different tasks. While these methods achieve comparable performance to RQT in single-task settings with manually chosen base models (e.g., CodeLLaMA-Python for code generation), their performance on code LLMs degrades catastrophically to near-zero levels as task diversity increases. As analyzed in Section 3.2, this performance collapse arises from the limited representation of extreme low-bit quantization when applied to large-magnitude model deltas, which result from large-sized fine-tuned datasets. In contrast, RQT avoids the explicit introduction of model deltas Δ by automatically allocating shared binary values across similar models, making it less sensitive to task diversity.

Table 1: Results of RQT for 16 models across four tasks, averaged over two evaluation datasets. Our method employs per-channel binary quantization for all nodes, except the 4-bit root node. This yields an average bit width of $(4 + 30)/16 = 2.125$. Previous PTQ methods, including GPTQ, QuaRot, and OmniQuant, utilize group-wise quantization with a group size of 128. For delta compression, we implement BitDelta and GPT-Zip, using the LLaMA-2-7B (16-bit floating-point) as the base model and quantizing 16 model deltas to 1-bit and 2-bit precision, respectively.

Method	#Bits	Math				Code			
		WizardMath	MetaMath	Xwin-Math	MuggleMath	CodeLLaMA-PY	Wizardcoder	Magocoder-S-CL	ReflectionCoder
FP	16	32.75	43.76	64.48	46.00	42.32	56.67	69.55	73.05
GPTQ	4.125	30.65	43.35	62.54	43.90	38.68	51.59	70.40	70.32
	3.125	21.00	38.85	56.02	34.35	23.48	38.61	55.55	57.54
	2.125	2.90	1.37	0.40	2.15	0.00	0.00	0.00	0.00
QuaRot	3.125	23.85	39.25	56.50	37.90	32.14	42.96	61.00	63.73
	2.125	0.90	1.85	1.25	1.50	0.00	0.00	0.00	0.00
OmniQuant	3.125	22.65	40.70	57.84	36.85	30.84	46.28	62.70	62.98
	2.125	1.95	8.75	2.86	4.95	4.15	1.61	17.70	18.73
BitDelta	2	30.95	41.05	57.57	42.65	0.00	0.00	0.00	0.00
GPT-Zip	3	32.10	43.40	63.21	45.70	0.00	0.00	0.00	0.00
RQT	2.125	35.45	43.35	61.97	48.45	31.57	51.19	64.05	62.43

		Chat				Chinese			
		LLaMA-2-Chat	Vicuna-v1.5	Vicuna-v1.5-16k	LLaMA-2	CN-Alpaca-2	CN-LLaMA-2	CN-Alpaca-2-16k	CN-LLaMA-2-16k
FP	16	74.69	74.72	75.25	76.18	40.97	28.06	37.54	27.49
GPTQ	4.125	73.73	74.27	74.67	75.73	40.05	26.91	36.56	25.84
	3.125	69.37	70.95	69.60	72.19	34.42	26.28	34.53	24.30
	2.125	26.26	32.40	28.62	31.37	24.43	24.21	25.73	26.84
QuaRot	3.125	69.99	71.50	66.81	69.65	32.91	26.22	34.52	25.97
	2.125	29.06	33.99	28.82	31.02	25.25	24.90	25.11	24.04
OmniQuant	3.125	70.50	71.47	70.92	73.31	33.79	26.33	33.73	26.31
	2.125	43.11	50.45	37.65	52.18	24.33	25.27	24.72	24.18
BitDelta	2	74.70	75.10	63.20	75.32	36.91	26.47	29.13	25.31
GPT-Zip	3	75.06	75.15	75.62	76.17	39.83	26.92	35.56	26.13
RQT	2.125	74.92	75.40	74.79	75.28	40.38	27.63	36.07	27.20

Methods	Math	Code	Chat	Chinese
Top-down single-step method	40.29	49.19	74.77	30.41
Our two-step method	42.99	50.74	74.87	32.91
– w/o model-specific scales	43.32	30.41	74.46	31.87
+ with scale distillation	47.30	52.58	75.09	32.76

Table 2: Design choices of building an RQT.

Models	Dataset	w2g128	w3g128	RQT
MAmmoTH-7B	MATH	1.38	21.94	25.44
LlaVA-v1.5	TextVQA	0.44	53.74	56.64
Magocoder-CL	MBPP	0.00	58.5.0	64.00
XwinLM-v0.2	HellaSwag	29.63	72.66	76.17

Table 3: Results of newly added models.

4.3 Ablation Studies

Metric	Math	Code	Chat	Chinese
L_2 norm	42.99	50.74	74.87	32.91
Cosine Similarity	43.56	50.86	74.95	31.78
Sign Difference	42.59	51.05	74.80	31.78

Table 4: Comparison of different distance metrics in tree construction.

In this section, we experiment with a single-step method in a top-down manner, which serves as an alternative to the two-step greedy approach described in Section 3.2. Specifically, at each node, we first compute the optimal quantization parameters using Equation (7) and Equation (6), and then

determine how to distribute the residual weights among the child nodes using Equation 9. More details can be found in the appendix.

Table 2 shows that the single-step method performs worse than the two-step approach introduced in Section 3.2. We believe this is because, in the single-step method, even negligible quantization errors in the current node can lead to different model assignments, making the algorithm less robust.

Table 2 also verifies the effectiveness of utilizing model-specific scales and scale distillation. Additionally, Table 4 investigates the impact of the distance metric on tree construction. Considering the similar performance, we ultimately chose the L_2 distance, consistent with the objective of resid-

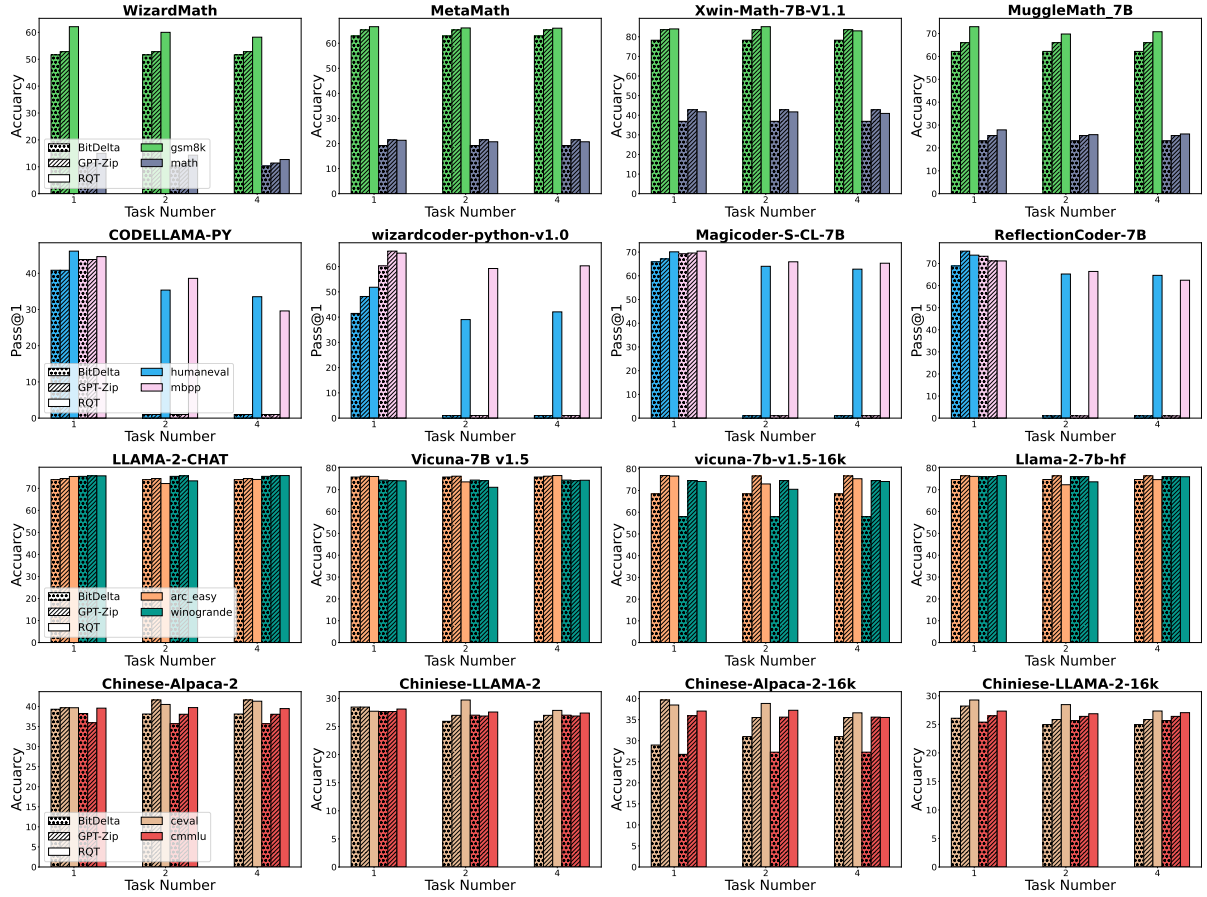


Figure 5: Results of RQT on different task numbers, where each task contains four models. Each color represents an evaluation dataset, and different hatch patterns indicate the various methods used. While delta compression methods achieve performance comparable to RQT when all models belong to a single task (the first six bars in each figure), their performance degrades significantly (to near-zero levels) on code LLMs when applied across multiple tasks. In contrast, RQT is less sensitive to delta magnitude and demonstrates greater robustness to task diversity.

ual quantization.

4.4 Incorporating New Models into the RQT

In this section, we describe the process of incorporating new models into an existing RQT without the need to rebuild the entire tree. This method also relaxes the constraint that the number of models, N , must be a power of 2. Specifically, we begin by identifying the optimal position to add the new leaf and then perform residual quantization from the root to the leaf. During this process, we only compute the optimal scales for the new models, without altering the shared low-bit values. The complete pseudocode for this process is provided in the appendix. We evaluate the addition of new models using MAMmoTH-7B (Yue et al., 2024), Llava-1.5 (Liu et al., 2023a) (evaluated on the TextVQA dataset (Singh et al., 2019)), Magicoder-CL (Wei et al., 2024), and Xwin-LM-v0.2 (Cui et al., 2023), with the existing RQT built upon the previous 16

models introduced in Section 4.1. The results, presented in Table 3, compare the performance of the newly added models with that of the GPTQ methods. In future work, we will explore dynamic adjustments to the RQT and develop more general construction methods.

5 Conclusion

In this paper, we introduce the Residual Quantization Tree (RQT), which mitigates the sensitivity of previous delta compression methods to the magnitude of delta values. We also present an efficient two-step greedy approach that separates model assignment from residual quantization. Experimental results demonstrate the effectiveness of the RQT when compared to traditional PTQ methods and delta-compression methods. Even with an average bitwidth of approximately 2 bits, RQT achieves only a 3% average accuracy degradation on fine-tuned LLaMA-2-7B models across a range of tasks.

Limitations

In this paper, we introduce the Residual Quantization Tree (RQT), a hierarchical quantization framework that automatically shares low-bit integer weights across similar fine-tuned models. While the RQT can, in principle, adopt various tree structures, the algorithm proposed in Section 3.2 is limited to a complete binary tree. Further research is needed to explore dynamic adjustments to the RQT and to develop a more general construction method. Additionally, although the RQT is primarily designed for multi-model compression, it also holds potential for applications in other contexts, such as cross-layer quantization.

References

- Saleh Ashkboos, Amirkeivan Mohtashami, Maximilian L. Croci, Bo Li, Martin Jaggi, Dan Alistarh, Torsten Hoefler, and James Hensman. 2024. [Quarot: Outlier-free 4-bit inference in rotated llms](#). *CoRR*, abs/2404.00456.
- Jacob Austin, Augustus Odena, Maxwell I. Nye, Maarten Bosma, Henryk Michalewski, David Dohan, Ellen Jiang, Carrie J. Cai, Michael Terry, Quoc V. Le, and Charles Sutton. 2021. [Program synthesis with large language models](#). *CoRR*, abs/2108.07732.
- Dan Biderman, Jose Javier Gonzalez Ortiz, Jacob P. Portes, Mansheej Paul, Philip Greengard, Connor Jennings, Daniel King, Sam Havens, Vitaliy Chiley, Jonathan Frankle, Cody Blakeney, and John P. Cunningham. 2024. [Lora learns less and forgets less](#). *CoRR*, abs/2405.09673.
- Lequn Chen, Zihao Ye, Yongji Wu, Danyang Zhuo, Luis Ceze, and Arvind Krishnamurthy. 2024. [Punica: Multi-tenant lora serving](#). In *Proceedings of the Seventeenth Annual Conference on Machine Learning and Systems, MLSys 2024, Santa Clara, CA, USA, May 13-16, 2024*. mlsys.org.
- Mark Chen, Jerry Tworek, Heewoo Jun, Qiming Yuan, Henrique Pondé de Oliveira Pinto, Jared Kaplan, Harri Edwards, Yuri Burda, Nicholas Joseph, Greg Brockman, Alex Ray, Raul Puri, Gretchen Krueger, Michael Petrov, Heidy Khlaaf, Girish Sastry, Pamela Mishkin, Brooke Chan, Scott Gray, Nick Ryder, Mikhail Pavlov, Alethea Power, Lukasz Kaiser, Mohammad Bavarian, Clemens Winter, Philippe Tillet, Felipe Petroski Such, Dave Cummings, Matthias Plappert, Fotios Chantzis, Elizabeth Barnes, Ariel Herbert-Voss, William Hebgen Guss, Alex Nichol, Alex Paino, Nikolas Tezak, Jie Tang, Igor Babuschkin, Suchir Balaji, Shantanu Jain, William Saunders, Christopher Hesse, Andrew N. Carr, Jan Leike, Joshua Achiam, Vedant Misra, Evan Morikawa, Alec Radford, Matthew Knight, Miles Brundage, Mira Murati, Katie Mayer, Peter Welinder, Bob McGrew, Dario Amodei, Sam McCandlish, Ilya Sutskever, and Wojciech Zaremba. 2021. [Evaluating large language models trained on code](#). *CoRR*, abs/2107.03374.
- Wei-Lin Chiang, Zhuohan Li, Zi Lin, Ying Sheng, Zhanghao Wu, Hao Zhang, Lianmin Zheng, Siyuan Zhuang, Yonghao Zhuang, Joseph E. Gonzalez, Ion Stoica, and Eric P. Xing. 2023. [Vicuna: An open-source chatbot impressing gpt-4 with 90%* chatgpt quality](#).
- Peter Clark, Isaac Cowhey, Oren Etzioni, Tushar Khot, Ashish Sabharwal, Carissa Schoenick, and Oyvind Tafjord. 2018. [Think you have solved question answering? try arc, the AI2 reasoning challenge](#). *CoRR*, abs/1803.05457.
- Karl Cobbe, Vineet Kosaraju, Mohammad Bavarian, Mark Chen, Heewoo Jun, Lukasz Kaiser, Matthias Plappert, Jerry Tworek, Jacob Hilton, Reiichiro Nakano, Christopher Hesse, and John Schulman. 2021. [Training verifiers to solve math word problems](#). *CoRR*, abs/2110.14168.
- Thomas H Cormen, Charles E Leiserson, Ronald L Rivest, and Clifford Stein. 2022. *Introduction to algorithms*. MIT press.
- Matthieu Courbariaux and Yoshua Bengio. 2016. [Binarynet: Training deep neural networks with weights and activations constrained to +1 or -1](#). *CoRR*, abs/1602.02830.
- Yiming Cui, Ziqing Yang, and Xin Yao. 2023. [Efficient and effective text encoding for chinese llama and alpaca](#). *arXiv preprint arXiv:2304.08177*.
- Ning Ding, Yujia Qin, Guang Yang, Fuchao Wei, Zonghan Yang, Yusheng Su, Shengding Hu, Yulin Chen, Chi-Min Chan, Weize Chen, Jing Yi, Weilin Zhao, Xiaozhi Wang, Zhiyuan Liu, Hai-Tao Zheng, Jianfei Chen, Yang Liu, Jie Tang, Juanzi Li, and Maosong Sun. 2023. [Parameter-efficient fine-tuning of large-scale pre-trained language models](#). *Nat. Mac. Intell.*, 5(3):220–235.
- Clémentine Fourrier, Nathan Habib, Alina Lozovskaya, Konrad Szafer, and Thomas Wolf. 2024. Open llm leaderboard v2. https://huggingface.co/spaces/open-llm-leaderboard/open_llm_leaderboard.
- Elias Frantar, Saleh Ashkboos, Torsten Hoefler, and Dan Alistarh. 2022. [GPTQ: accurate post-training quantization for generative pre-trained transformers](#). *CoRR*, abs/2210.17323.
- Google. 2023. [Vertex AI](#).
- Dan Hendrycks, Collin Burns, Saurav Kadavath, Akul Arora, Steven Basart, Eric Tang, Dawn Song, and Jacob Steinhardt. 2021. [Measuring mathematical problem solving with the MATH dataset](#). In *Proceedings of the Neural Information Processing Systems Track on Datasets and Benchmarks 1, NeurIPS Datasets and Benchmarks 2021, December 2021, virtual*.

621	Yuzhen Huang, Yuzhuo Bai, Zhihao Zhu, Junlei	Shi, Raghuraman Krishnamoorthi, and Vikas Chan-	676
622	Zhang, Jinghan Zhang, Tangjun Su, Junteng Liu,	dra. 2023b. LLM-QAT: data-free quantization	677
623	Chuancheng Lv, Yikai Zhang, Jiayi Lei, Yao	aware training for large language models . <i>CoRR</i> ,	678
624	Fu, Maosong Sun, and Junxian He. 2023. C-	abs/2305.17888.	679
625	eval: A multi-level multi-discipline chinese evalu-		
626	ation suite for foundation models. <i>arXiv preprint</i>	Zechun Liu, Changsheng Zhao, Igor Fedorov, Bilge	680
627	<i>arXiv:2305.08322</i> .	Soran, Dhruv Choudhary, Raghuraman Krishnamoor-	681
628	Berivan Isik, Hermann Kumbong, Wanyi Ning, Xiaozhe	thi, Vikas Chandra, Yuandong Tian, and Tijmen	682
629	Yao, Sanmi Koyejo, and Ce Zhang. 2023. Gpt-zip:	Blankevoort. 2024b. Spinquant: LLM quantization	683
630	Deep compression of finetuned large language mod-	with learned rotations . <i>CoRR</i> , abs/2405.16406.	684
631	els. In <i>Workshop on Efficient Systems for Foundation</i>		
632	<i>Models@ ICML2023</i> .	Haipeng Luo, Qingfeng Sun, Can Xu, Pu Zhao, Jian-	685
633	Benoit Jacob, Skirmantas Kligys, Bo Chen, Menglong	guang Lou, Chongyang Tao, Xiubo Geng, Qingwei	686
634	Zhu, Matthew Tang, Andrew G. Howard, Hartwig	Lin, Shifeng Chen, and Dongmei Zhang. 2023. Wiz-	687
635	Adam, and Dmitry Kalenichenko. 2018. Quantiza-	ardmath: Empowering mathematical reasoning for	688
636	tion and training of neural networks for efficient	large language models via reinforced evol-instruct .	689
637	integer-arithmetic-only inference . In <i>2018 IEEE Con-</i>	<i>CoRR</i> , abs/2308.09583.	690
638	<i>ference on Computer Vision and Pattern Recognition,</i>		
639	<i>CVPR 2018, Salt Lake City, UT, USA, June 18-22,</i>	Ziyang Luo, Can Xu, Pu Zhao, Qingfeng Sun, Xi-	691
640	<i>2018</i> , pages 2704–2713. Computer Vision Founda-	ubo Geng, Wenxiang Hu, Chongyang Tao, Jing Ma,	692
641	tion / IEEE Computer Society.	Qingwei Lin, and Daxin Jiang. 2024. Wizardcoder:	693
642	Chen Li, Weiqi Wang, Jingcheng Hu, Yixuan Wei, Nan-	Empowering code large language models with evol-	694
643	ning Zheng, Han Hu, Zheng Zhang, and Houwen	instruct . In <i>The Twelfth International Conference</i>	695
644	Peng. 2024. Common 7b language models al-	<i>on Learning Representations, ICLR 2024, Vienna,</i>	696
645	ready possess strong math capabilities . <i>CoRR</i> ,	<i>Austria, May 7-11, 2024</i> . OpenReview.net.	697
646	abs/2403.04706.		
647	Chengpeng Li, Zheng Yuan, Hongyi Yuan, Guanting	Microsoft. 2023. Fine-tune a foundation model from the	698
648	Dong, Keming Lu, Jiancan Wu, Chuanqi Tan, Xiang	model catalog in Azure Machine Learning - Training .	699
649	Wang, and Chang Zhou. 2023a. Query and response		
650	augmentation cannot help out-of-domain math rea-	Markus Nagel, Rana Ali Amjad, Mart van Baalen, Chris-	700
651	soning generalization . <i>CoRR</i> , abs/2310.05506.	tos Louizos, and Tijmen Blankevoort. 2020. Up or	701
652	Haonan Li, Yixuan Zhang, Fajri Koto, Yifei Yang, Hai	down? adaptive rounding for post-training quantiza-	702
653	Zhao, Yeyun Gong, Nan Duan, and Timothy Bald-	tion . In <i>Proceedings of the 37th International Con-</i>	703
654	win. 2023b. Cmmlu: Measuring massive multi-	<i>ference on Machine Learning, ICML 2020, 13-18</i>	704
655	task language understanding in chinese . <i>Preprint,</i>	<i>July 2020, Virtual Event</i> , volume 119 of <i>Proceedings</i>	705
656	arXiv:2306.09212.	<i>of Machine Learning Research</i> , pages 7197–7206.	706
657	Yuhang Li, Ruihao Gong, Xu Tan, Yang Yang, Peng	PMLR.	707
658	Hu, Qi Zhang, Fengwei Yu, Wei Wang, and Shi Gu.		
659	2021. BRECQ: pushing the limit of post-training	OpenAI. 2023a. GPT-3.5 Turbo fine-tuning and API	708
660	quantization by block reconstruction . In <i>9th Inter-</i>	updates .	709
661	<i>national Conference on Learning Representations,</i>		
662	<i>ICLR 2021, Virtual Event, Austria, May 3-7, 2021</i> .	OpenAI. 2023b. GPT-4 technical report . <i>CoRR</i> ,	710
663	OpenReview.net.	abs/2303.08774.	711
664	Ji Lin, Jiaming Tang, Haotian Tang, Shang Yang,	Kuntal Kumar Pal, Kazuaki Kashihara, Ujjwala Anan-	712
665	Xingyu Dang, and Song Han. 2023. AWQ: activation-	theswaran, Kirby C. Kuznia, Siddhesh Jagtap, and	713
666	aware weight quantization for LLM compression and	Chitta Baral. 2023. Exploring the limits of trans-	714
667	acceleration . <i>CoRR</i> , abs/2306.00978.	fer learning with unified model in the cybersecurity	715
668	Haotian Liu, Chunyuan Li, Qingyang Wu, and Yong Jae	domain . <i>CoRR</i> , abs/2302.10346.	716
669	Lee. 2023a. Visual instruction tuning.		
670	James Liu, Guangxuan Xiao, Kai Li, Jason D. Lee,	Bowen Ping, Shuo Wang, Hanqing Wang, Xu Han,	717
671	Song Han, Tri Dao, and Tianle Cai. 2024a. Bitdelta:	Yuzhuang Xu, Yukun Yan, Yun Chen, Baobao	718
672	Your fine-tune may only be worth one bit . <i>CoRR</i> ,	Chang, Zhiyuan Liu, and Maosong Sun. 2024.	719
673	abs/2402.10193.	Delta-come: Training-free delta-compression with	720
674	Zechun Liu, Barlas Oguz, Changsheng Zhao, Ernie	mixed-precision for large language models . <i>CoRR</i> ,	721
675	Chang, Pierre Stock, Yashar Mehdad, Yangyang	abs/2406.08903.	722
		Mohammad Rastegari, Vicente Ordonez, Joseph Red-	723
		mon, and Ali Farhadi. 2016. Xnor-net: Imagenet	724
		classification using binary convolutional neural net-	725
		works . In <i>Computer Vision - ECCV 2016 - 14th</i>	726
		<i>European Conference, Amsterdam, The Netherlands,</i>	727
		<i>October 11-14, 2016, Proceedings, Part IV</i> , volume	728
		9908 of <i>Lecture Notes in Computer Science</i> , pages	729
		525–542. Springer.	730

731	Houxing Ren, Mingjie Zhan, Zhongyuan Wu, Aojun Zhou, Juntong Pan, and Hongsheng Li. 2024. Reflectioncoder: Learning from reflection sequence for enhanced one-off code generation . <i>Preprint</i> , arXiv:2405.17057.	788
732		789
733		790
734		791
735		
736	Baptiste Rozière, Jonas Gehring, Fabian Gloeckle, Sten Sootla, Itai Gat, Xiaoqing Ellen Tan, Yossi Adi, Jingyu Liu, Tal Remez, Jérémy Rapin, Artyom Kozhevnikov, Ivan Evtimov, Joanna Bitton, Manish Bhatt, Cristian Canton-Ferrer, Aaron Grattafiori, Wenhan Xiong, Alexandre Défossez, Jade Copet, Faisal Azhar, Hugo Touvron, Louis Martin, Nicolas Usunier, Thomas Scialom, and Gabriel Synnaeve. 2023. Code llama: Open foundation models for code . <i>CoRR</i> , abs/2308.12950.	792
737		793
738		794
739		795
740		
741		796
742		797
743		798
744		799
745		800
746		801
747	Simo Ryu, Seunghyun Seo, and Jaejun Yoo. 2023. Efficient storage of fine-tuned models via low-rank approximation of weight residuals . <i>CoRR</i> , abs/2305.18425.	802
748		
749		803
750		804
751	Wenqi Shao, Mengzhao Chen, Zhaoyang Zhang, Peng Xu, Lirui Zhao, Zhiqian Li, Kaipeng Zhang, Peng Gao, Yu Qiao, and Ping Luo. 2023. Omniquant: Omnidirectionally calibrated quantization for large language models . <i>CoRR</i> , abs/2308.13137.	805
752		
753		806
754		807
755		808
756	Ying Sheng, Shiyi Cao, Dacheng Li, Coleman Hooper, Nicholas Lee, Shuo Yang, Christopher Chou, Banghua Zhu, Lianmin Zheng, Kurt Keutzer, Joseph E. Gonzalez, and Ion Stoica. 2023. S-lora: Serving thousands of concurrent lora adapters . <i>CoRR</i> , abs/2311.03285.	809
757		810
758		811
759		
760		812
761	Amanpreet Singh, Vivek Natarjan, Meet Shah, Yu Jiang, Xinlei Chen, Devi Parikh, and Marcus Rohrbach. 2019. Towards vqa models that can read. In <i>Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition</i> , pages 8317–8326.	813
762		814
763		815
764		816
765		817
766		818
767	Hugo Touvron, Thibaut Lavril, Gautier Izacard, Xavier Martinet, Marie-Anne Lachaux, Timothée Lacroix, Baptiste Rozière, Naman Goyal, Eric Hambro, Faisal Azhar, Aurélien Rodriguez, Armand Joulin, Edouard Grave, and Guillaume Lample. 2023. Llama: Open and efficient foundation language models . <i>CoRR</i> , abs/2302.13971.	819
768		820
769		821
770		822
771		823
772		824
773		825
774	Xiuying Wei, Yunchen Zhang, Xiangguo Zhang, Ruihao Gong, Shanghang Zhang, Qi Zhang, Fengwei Yu, and Xianglong Liu. 2022. Outlier suppression: Pushing the limit of low-bit transformer language models . In <i>NeurIPS</i> .	826
775		
776		
777		
778	Yuxiang Wei, Zhe Wang, Jiawei Liu, Yifeng Ding, and Lingming Zhang. 2024. Magicoder: Empowering code generation with oss-instruct . In <i>Forty-first International Conference on Machine Learning, ICML 2024, Vienna, Austria, July 21-27, 2024</i> . OpenReview.net.	
779		
780		
781		
782		
783		
784	Guangxuan Xiao, Ji Lin, Mickaël Seznec, Hao Wu, Julien Demouth, and Song Han. 2023. Smoothquant: Accurate and efficient post-training quantization for large language models . In <i>International Conference on Machine Learning, ICML 2023, 23-29 July 2023, Honolulu, Hawaii, USA</i> , volume 202 of <i>Proceedings of Machine Learning Research</i> , pages 38087–38099. PMLR.	
785		
786		
787		

A Appendix

A.1 Detailed Experiment Results

In Section 4, due to page limitations, we present only the aggregated results from two evaluation datasets. The detailed results can be found in Table 5 to Table 8.

A.2 Visualization of the RQT

Figure 6 shows a visualization of the RQT applied to the weights of the down projection layer in the first decoder layer. We observe that (a) similar models are positioned closer together in the RQT and (b) the reconstruction errors for each model decrease as the residual quantization performed from top to down. These observations validate the effectiveness of our algorithm.

A.3 More details about ablation studies

The algorithm of top-down single-step method

Although we propose a two-step greedy approach for constructing an RQT, we also experimented with a single-step method in a top-down manner, as shown in Table 2. Specifically, at each node, we first optimize the quantization parameters using Equation (7) and Equation (6), and then determine how to distribute the residual weights among the child nodes by Equation (9)

$$\begin{aligned} \min_{\{x_i\}_{1 \leq i \leq N}} \quad & \sum_{1 \leq i < j \leq N} I(x_i = x_j) \cdot d_{ij}, \\ \text{subject to:} \quad & \sum_{1 \leq i \leq N} x_i = \frac{N}{2}, \quad x_i \in \{0, 1\}. \end{aligned} \quad (9)$$

Here, x_i is a binary variable indicating whether node i belongs to the first set. The indicator function $I(x_i = x_j)$ equals 1 when $x_i = x_j$ (W_i and W_j belong to the same set), and 0 otherwise. Additionally, d_{ij} represents the distance between the weight matrices W_i and W_j , and it is chosen as the L_2 distance:

$$d_{ij} = \|W_i - W_j\|_2. \quad (10)$$

The complete pseudocode is formally described in Algorithm 3.

The algorithm of model addition. In section 4.4, we discuss how to incorporate new models into an existing RQT without rebuilding the entire tree and relax the constraint that N is a power of 2. Specifically, we first search for the best position to add the new leaf and then perform residual

Algorithm 3 Single_Step_Method

```

1: Input: A set of weights  $\Theta = \{W_i \mid i = 1, \dots, N\}$ 
2: Output: A complete residual quantization tree  $T$  with defined  $Q$  and  $\Sigma$ 
3: while not converged do
4:   Compute  $\alpha_i$  using Equation 6
5:   Compute  $Q$  using Equation 7
6: end while
7: Update  $W_i \leftarrow W_i - Q \cdot \alpha_i$  for all  $i$ 
8: Create child nodes  $T_{\text{left}}, T_{\text{right}}$ 
9: Set  $C[T] \leftarrow \{T_{\text{left}}, T_{\text{right}}\}$ 
10: Compute  $S[T_{\text{right}}]$  and  $S[T_{\text{left}}]$  using Equation 9
11: for each  $T_j \in C[T]$  do
12:   Single_Step_Method( $\{W_i \mid i \in S[T_j]\}, T_j$ )
13: end for
14: return  $T$ 

```

Algorithm 4 Model_Addition

```

1: Input: A set of weights  $\Theta = \{W_i \mid i = 1, \dots, N\}$ 
2: Input: Newly added weight  $W_{N+1}$ 
3: Input: A residual quantization tree  $T$  build on  $\Theta$ 
4:  $k = \arg \min ||W_{N+1} - W_j||$  for  $j \in \{1, \dots, N\}$ 
5: Let  $T_k$  denote the leaf node where  $S[T_k] = \{k\}$ 
6: while true do
7:   Compute  $\alpha_{N+1}$  by Equation (6)
8:   Set  $\Sigma[T] \leftarrow \Sigma[T] \cup \{\alpha_{N+1}\}$ 
9:   Set  $S[T] \leftarrow S[T] \cup \{N+1\}$ 
10:  Set  $W_{N+1} \leftarrow W_{N+1} - \alpha_{N+1} \cdot Q[T]$ 
11:  if  $T = \text{parent}(T_k)$  then
12:    break
13:  else
14:     $T = T_j$  where  $T_j \in C[T]$  and  $k \in T_j$ 
15:  end if
16: end while
17: Create a new leaf node  $T_{\text{new}}$ 
18:  $S[T_{\text{new}}] = \{N+1\}$ 
19:  $C[\text{parent}(T_k)] = C[\text{parent}(T_k)] \cup \{T_{\text{new}}\}$ 
20: Compute  $Q[T_{\text{new}}]$  and  $\alpha_{N+1}[T_{\text{new}}]$ 
21: return  $T$ 

```

quantization from the root to the leaf, only computing optimal scales for the new models without changing the shared low-bit values. The complete pseudocode is formally described in Algorithm 4.

Table 5: The detailed results of experiments on math LLMs.

MODEL DATASET	WizardMath-7B			MetaMath-7B			Xwin-Math-7B-V1.1			MuggleMath-7B		
	GSM8K	MATH	Avg.	GSM8K	MATH	Avg.	GSM8K	MATH	Avg.	GSM8K	MATH	Avg.
FP	54.20	11.30	32.75	66.71	20.80	43.76	84.53	44.42	64.48	67.10	24.90	46.00
GPTQ-w4-g128	51.30	10.00	30.65	66.60	20.10	43.35	82.90	42.18	62.54	63.80	24.00	43.90
GPTQ-W3-g128	35.90	6.10	21.00	61.00	16.70	38.85	75.57	36.46	56.02	53.40	15.30	34.35
GPTQ-W2-g128	3.10	2.70	2.90	1.10	1.63	1.37	0.37	0.42	0.40	2.00	2.30	2.15
QuaRot-w2g128	1.30	0.50	0.90	2.70	1.00	1.85	0.90	1.60	1.25	1.50	1.1	1.50
QuaRot-w3g128	40.80	6.90	23.85	61.50	17.00	39.25	76.95	36.04	56.50	58.50	17.30	37.90
Omniquant-w3g128	38.50	6.80	22.65	63.50	17.90	40.70	79.30	36.38	57.84	55.00	18.70	36.85
Omniquant-w2g128	2.30	1.60	1.95	14.30	3.20	8.75	4.01	1.70	2.86	6.70	3.20	4.95
BitDelta (1-bit Delta)	51.60	10.30	30.95	62.90	19.20	41.05	78.24	36.90	57.57	62.20	23.10	42.65
GPT-Zip (2-bit Delta)	52.80	11.40	32.10	65.30	21.50	43.40	83.62	42.80	63.21	66.00	25.40	45.70
RQT	58.20	12.70	35.45	66.00	20.70	43.35	83.01	40.92	61.97	70.80	26.10	48.45
RQT (w/o scales distillation)	57.60	13.10	35.35	65.40	20.50	42.95	82.94	40.90	61.92	69.80	25.00	47.40
RQT (w/o model-specific scales)	59.10	12.80	35.95	66.20	20.40	28.30	81.12	40.16	60.64	70.80	26.00	48.40
Top-down Single-step method	49.70	10.10	29.90	65.40	20.20	29.30	78.92	36.06	57.49	65.70	23.20	44.45
Cosine Similarity	57.50	13.00	35.25	65.00	21.30	30.30	83.01	41.88	62.45	68.70	23.80	46.25
Sign difference	47.90	10.10	29.00	66.60	20.30	31.30	83.26	41.52	62.39	70.40	24.90	47.65

Table 6: The detailed experimental results on code LLMs are provided. HE refers to the HumanEval (Chen et al., 2021) dataset.

MODEL DATASET	CodeLLaMA-Python-7B			WizardCoder-7B			Magicoder-S-CL-7B			ReflectionCoder-7B		
	HE	MBPP	Avg.	HE	MBPP	Avg.	HE	MBPP	Avg.	HE	MBPP	Avg.
FP	40.84	43.80	42.32	48.78	64.55	56.67	69.50	69.60	69.55	74.40	71.69	73.05
GPTQ-w4-g128	34.75	42.60	38.68	42.07	61.11	51.59	72.00	68.80	70.40	73.17	67.46	70.32
GPTQ-W3-g128	22.56	24.40	23.48	29.87	47.35	38.61	51.80	59.30	55.55	56.09	58.99	57.54
GPTQ-W2-g128	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00
QuaRot-w2g128	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00
QuaRot-w3g128	29.87	34.40	32.14	33.53	52.38	42.96	60.40	61.60	61.00	61.58	65.87	63.73
Omniquant-w3g128	31.20	30.48	30.84	37.80	54.76	46.28	64.00	61.40	62.70	60.36	65.60	62.98
Omniquant-w2g128	6.09	2.20	4.15	2.43	0.79	1.61	7.90	27.50	17.70	15.24	22.22	18.73
BitDelta (1-bit Delta)	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00
GPT-Zip (2-bit Delta)	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00
RQT	33.53	29.60	31.57	42.07	60.31	51.19	62.80	65.30	64.05	64.63	62.43	62.43
RQT (w/o scales distillation)	35.36	33.60	34.48	38.41	57.14	47.78	54.30	64.00	59.15	60.97	62.16	61.57
RQT (w/o model-specific scales)	43.32	15.85	20.80	18.33	7.31	34.65	20.98	27.40	46.80	37.10	37.80	52.64
Top-down Single-step method	40.29	30.60	34.14	32.37	34.14	56.08	45.11	54.90	60.80	57.85	60.97	61.90
Cosine Similarity	43.56	32.31	32.40	32.36	38.41	55.82	47.12	58.50	63.80	61.15	61.58	64.02
Sign difference	42.59	32.31	32.60	32.46	38.41	56.08	47.25	58.50	64.00	61.25	62.19	64.28

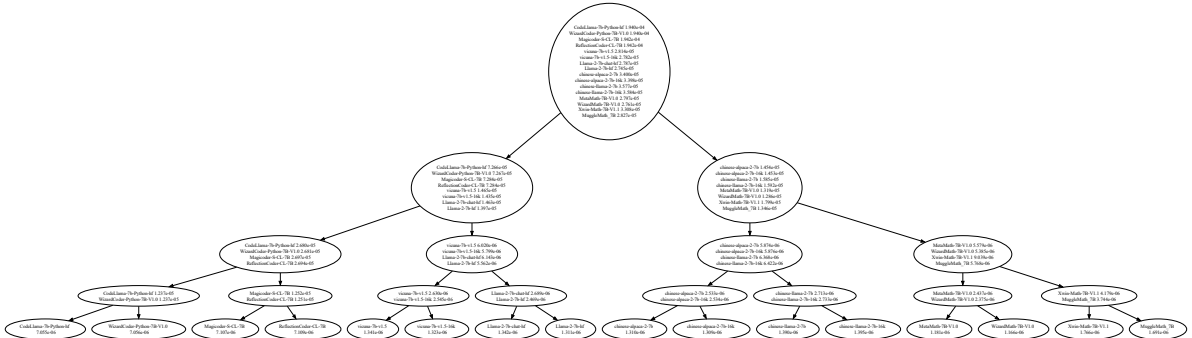


Figure 6: Visualization of the Residual Quantization Tree (RQT) applied to the down projector of the first decoder layer, with the MSE of each model annotated.

Table 7: The detailed results of experiments on chat LLMs are presented. ARC-E refers to the ARC-Easy (Clark et al., 2018) dataset, and HS refers to the HellaSwag (Zellers et al., 2019) dataset.

MODEL DATASET	LLaMA-2-chat-7B			Vicuna-7B v1.5			Vicuna-7B-v1.5-16k			LLaMA-2-7B		
	ARC-E	HS	Avg.	ARC-E	HS	Avg.	ARC-E	HS	Avg.	ARC-E	HS	Avg.
FP	73.91	75.47	74.69	75.63	73.81	74.72	76.18	74.31	75.25	76.35	76.00	76.18
GPTQ-w4-g128	72.77	74.69	73.73	75.34	73.20	74.27	76.26	73.07	74.67	75.84	75.62	75.73
GPTQ-W3-g128	67.30	71.44	69.37	72.22	69.68	70.95	69.99	69.20	69.60	72.31	72.06	72.19
GPTQ-W2-g128	25.29	27.23	26.26	32.66	32.14	32.40	29.59	27.65	28.62	31.52	31.22	31.37
QuaRot-w2g128	29.04	29.08	29.06	35.73	32.24	33.99	29.67	27.96	28.82	31.48	30.55	31.02
QuaRot-w3g128	68.81	71.17	69.99	73.02	69.98	71.50	69.87	63.74	66.81	69.44	69.85	69.65
Omniquant-w3g128	69.57	71.43	70.50	72.64	70.30	71.47	72.58	69.25	70.92	74.07	72.54	73.31
Omniquant-w2g128	43.48	42.73	43.11	53.75	47.14	50.45	37.37	37.92	37.65	52.95	51.40	52.18
BitDelta (1-bit Delta)	73.95	75.44	74.70	75.80	74.40	75.10	68.48	57.92	63.20	74.66	75.98	75.32
GPT-Zip (2-bit Delta)	74.45	75.66	75.06	76.18	74.11	75.15	76.68	74.56	75.62	76.35	75.98	76.17
RQT	74.03	75.80	74.92	76.47	74.32	75.40	75.42	74.15	74.79	74.62	75.93	75.28
RQT (w/o scales distillation)	74.33	74.98	74.66	76.22	73.70	74.96	75.04	73.57	74.31	75.04	76.08	75.56
RQT (w/o model-specific scales)	45.22	30.41	37.86	74.20	74.03	75.88	73.04	74.46	74.96	72.60	73.78	75.93
Top-down Single-step method	61.44	49.19	74.83	74.78	74.81	75.38	73.17	74.28	76.09	72.89	74.49	75.55
Cosine Similarity	62.80	50.86	75.00	74.92	74.96	75.97	73.47	74.72	75.08	73.61	74.35	75.76
Sign difference	63.24	51.05	74.33	74.16	74.25	75.88	73.56	74.72	75.13	73.63	74.38	75.63

Table 8: The detailed results of experiments on Chinese LLMs.

MODEL DATASET	Chinese-Alpaca-2			Chinese-LLAMA-2			Chinese-Alpaca-2-16k			Chinese-LLAMA-2-16k		
	CEval	CMMLU	Avg.	CEval	CMMLU	Avg.	CEval	CMMLU	Avg.	CEval	CMMLU	Avg.
FP	42.05	39.88	40.97	28.45	27.66	28.06	37.89	37.18	37.54	28.45	26.52	27.49
GPTQ-w4-g128	41.08	39.02	40.05	26.52	27.29	26.91	37.41	35.71	36.56	25.92	25.75	25.84
GPTQ-W3-g128	34.99	33.84	34.42	26.52	26.04	26.28	35.36	33.69	34.53	23.32	25.28	24.30
GPTQ-W2-g128	24.14	24.71	24.43	23.84	24.57	24.21	26.37	25.09	25.73	28.15	25.52	26.84
QuaRot-w2g128	25.70	24.80	25.25	25.85	24.90	24.90	25.18	25.03	25.11	22.88	25.19	24.04
QuaRot-w3g128	32.76	33.05	32.91	26.59	25.84	26.22	36.55	32.49	34.52	26.96	24.98	25.97
Omniquant-w3g128	35.14	32.43	33.79	26.82	25.83	26.33	35.21	32.24	33.73	26.74	25.87	26.31
Omniquant-w2g128	23.25	25.41	24.33	23.03	25.27	25.27	23.47	24.72	24.72	22.95	25.40	24.18
BitDelta (1-bit Delta)	38.11	35.70	36.91	25.92	27.01	26.47	30.98	27.28	29.13	24.96	25.66	25.31
GPT-Zip (2-bit Delta)	41.60	38.05	39.83	26.98	26.86	26.92	35.51	35.60	35.56	25.85	26.40	26.13
RQT	41.30	39.46	40.38	27.86	27.39	27.63	36.62	35.52	36.07	27.34	27.05	27.20
RQT (w/o scales distillation)	40.71	37.94	39.33	28.52	28.58	28.55	37.29	35.10	36.20	27.56	27.54	27.55
RQT (w/o model-specific scales)	75.17	75.55	74.46	40.56	37.30	38.93	27.34	27.14	27.24	36.47	34.53	35.50
Top-down Single-step method	75.50	75.53	74.77	35.58	35.35	35.47	26.96	26.54	26.75	34.99	32.99	33.99
Cosine Similarity	75.80	75.78	74.95	38.03	36.04	37.04	28.08	28.03	28.06	36.47	34.38	35.43
Sign difference	76.09	75.86	74.80	38.03	36.04	37.04	28.08	28.03	28.06	36.47	34.38	35.43