

CAT: POST-TRAINING QUANTIZATION ERROR REDUCTION VIA CLUSTER-BASED AFFINE TRANSFORMATION

Ali Zoljodi* & Masoud Daneshtalab

Department of Intelligent Future Technologies (IFT)
Mälardalen University
Västerås, 72123, Sweden
{ali.zoljodi, masoud.daneshtalab}@mdu.se

Radu Timofte

Computer Vision Lab, CAIDAS & IFI
University of Würzburg
John Skilton Str. 4a, 97074 Würzburg, Germany
{radu.timofte}@uni-wuerzburg.de

ABSTRACT

Post-Training Quantization (PTQ) reduces the memory footprint and computational overhead of deep neural networks by converting full-precision (FP) values into quantized and compressed data types. While PTQ is more cost-efficient than Quantization-Aware Training (QAT), it is highly susceptible to accuracy degradation under a low-bit quantization (LQ) regime (e.g., 2-bit). Affine transformation is a classical technique used to reduce the discrepancy between the information processed by a quantized model and that processed by its full-precision counterpart; however, we find that using plain affine transformation, which applies a uniform affine parameter set for all outputs, worsens the results in low-bit PTQ. To address this, we propose Cluster-based Affine Transformation (CAT), an error-reduction framework that employs cluster-specific parameters to align LQ outputs with FP counterparts. CAT refines LQ outputs with only a negligible number of additional parameters, without requiring fine-tuning of the model or quantization parameters. We further introduce a novel PTQ framework integrated with CAT. Experiments on ImageNet-1K show that this framework consistently outperforms prior PTQ methods across diverse architectures and LQ settings, achieving up to 53.18% Top-1 accuracy on W2A2 ResNet-18. Moreover, CAT enhances existing PTQ baselines by more than 3% when used as a plug-in. We plan to release our implementation alongside the publication of this paper.

1 INTRODUCTION

Deep neural networks (DNNs) achieve remarkable performance on different computer vision tasks Ravi et al. (2025); Xiao et al. (2025); Zhu et al. (2025); Li et al. (2025b); Wang et al. (2025); Zanella et al. (2024). but demand prohibitive memory and computation due to millions of floating-point parameters. Quantization, which reduces the precision of weights and activations to low bit-widths, offers an efficient compression strategy and is now widely supported by modern hardware. Post-Training Quantization (PTQ) is a practical model compression which determines quantization parameters without retraining or fine-tuning, requiring only a small calibration set. Despite this promise, PTQ suffers from a severe accuracy degradation when pushed to LQ regimes. In this work, we investigate the affine transformation Weisstein (2004); Ma et al. (2024) ability to reduce PTQ output errors to refine its accuracy degradation.

*Corresponding author

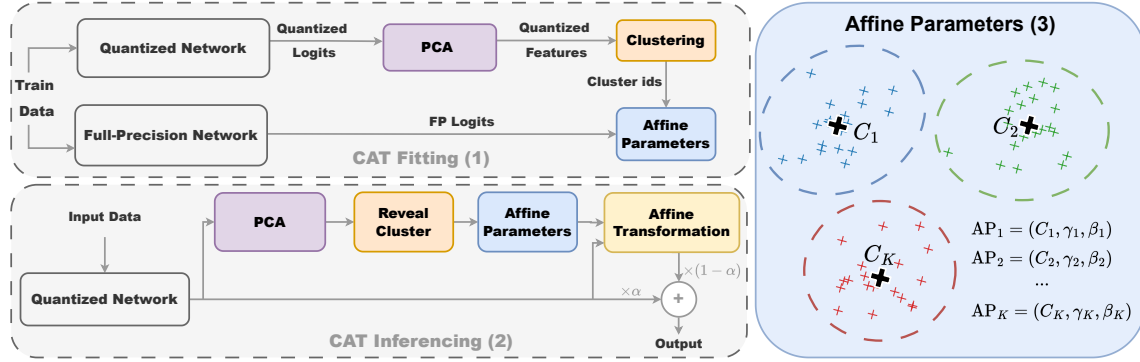


Figure 1: Cluster-based Affine Transformation (CAT); (1): CAT Fitting is the process of fitting the clustering model and estimating γ and β parameters for each cluster.(2) CAT Inferencing, the process of error reduction using CAT. (3) Affine Parameters (AP), a set of parameters we store for each cluster of CAT.

We first investigate a uniform set of affine transformation parameters, referred to as a plain affine transformation. As shown in Table 1, this approach fails to recover predictions and often further degrades the Top-1 accuracy of LQ ResNet-18. To alleviate the issue, we leverage clustering PTQ outputs and find a specific affine transformation parameter set for each cluster. Our proposed method, Cluster-based Affine Transformation (CAT), yields superior error reduction results for PTQ and especially the low-bit quantization (LQ) (e.g., 2-bit) regime. This strategy substantially reduces the output gap between FP and LQ, improving accuracy under low-bit settings with only a negligible number of additional affine parameters. Building on the advantages of CAT, we propose a novel PTQ framework. Compared to the baseline, our proposed PTQ framework raises the Top-1 accuracy of 2-bit quantized (W2A2) ResNet-18 to 53.18%. For compact DNNs such as MNasX2, CAT achieves a +1% accuracy improvement in the 2-bit quantization of both weight (W) and activation (A) (W2A2). Because of its ability to directly reduce output errors, CAT can be used as a plug-in for a wide range of PTQ methods. Our results show that, with this capability, CAT improves Top-1 accuracy by more than 3% for some PTQ methods. Our contributions can be summarized as follows:

- We propose Cluster-Affine Transformation (CAT), a novel output-level error reduction method that leverages the natural clusterability of logits to improve alignment.
- We introduce a novel state-of-the-art post-training quantization framework, which achieves consistent accuracy improvements across diverse architectures and LQ settings, surpassing prior PTQ methods with negligible overhead.

w	a	Method	Acc (%)
2	2	No affine transformation	52.84
		Plain affine transformation	52.32
		CAT (Ours)	53.18
4	2	No affine transformation	58.58
		Plain affine transformation	58.22
		CAT (Ours)	58.80
2	4	No affine transformation	65.12
		Plain affine transformation	65.17
		CAT (Ours)	65.25
4	4	No affine transformation	69.17
		Plain affine transformation	69.14
		CAT (Ours)	69.27

Table 1: Top-1 accuracy of ResNet-18 under quantization: comparison between no affine transformation, plain affine transformation, and CAT (ours).

- We achieve higher Top-1 accuracy compare to PTQ baselines on ImageNet-1k for different DNNs such as ResNet-18/50, MobileNetV2, and RegNetX.

2 RELATED WORK

Quantization Li et al. (2023); Sun et al. (2022); Qin et al. (2025); Cai et al. (2020); Harma et al. (2025); Li et al. (2025a; 2024); Zhou et al. (2025); Saxena et al. (2025) is a widely used model compression technique for deep neural networks (DNNs) that reduces model size and accelerates inference by representing weights and activations with lower bit precision. In practice, two paradigms exist: quantization-aware training (QAT), which incorporates quantization during model re-training (achieving high accuracy but at the cost of additional training on full datasets), and post-training quantization (PTQ), which converts a pre-trained model to low-bit format using only a small unlabeled calibration set without full re-training. While QAT preserves accuracy better, PTQ is efficient in development.

Post-training Quantization (PTQ) Banner et al. (2019); Liu et al. (2023); Nahshan et al. (2021); Banner et al. (2019); Nagel et al. (2020); Wang et al. (2020); Wei et al. (2022); Yuan et al. (2022); Lin et al. (2021); Ding et al. (2022); Lee et al. (2024); Shi et al. (2025); Ding et al. (2025); Zhong et al. (2025); Gong et al. (2025); Wu et al. (2025); Chen et al. (2025); Shen et al. (2025) is the process of determining quantization scale factors and zero-point without retraining or fine-tuning a model’s weights. The key challenge in PTQ is estimating the minimum and maximum values of weights and activations, and identifying outliers to clip. Early works on low-bit post-training quantization employed analytic methods to derive optimal clipping thresholds. Banner et al. (2019) proposed limiting activation ranges by statistically deriving activation distributions of tensors and determining the per-channel bit-width. Recent PTQ research has introduced methods to minimize accuracy loss by optimizing layer-wise or block-wise reconstructions. AdaRound Nagel et al. (2020) is a layer-wise reconstruction method that minimizes the local loss by adapting scale factors. To cover cross-layer interaction, BRECQ Li et al. (2021) extends layer-wise reconstruction to a group of layers (blocks) with second-order error approximations. PD-Quant Liu et al. (2023) proposes reconstructing layers and blocks via global loss minimization by comparing the network outputs before and after quantization. To improve the stability of PTQ, QDrop Wei et al. (2022) randomly drops activation quantization during calibration to improve generalization and robustness. Some other works adopt different PTQ formulations. Mr.BiQ Jeon et al. (2022) introduces a non-uniform, multi-level binary quantizer, where both scaling factors and binary codes are treated as learnable parameters and optimized jointly to minimize block-wise reconstruction error. Prepositive Feature Quantization (PFQ) Chu et al. (2024) reorganizes the PTQ framework by moving feature quantization before, rather than after, each layer. In practice, PFQ requires different calibration schedules to effectively align quantized and FP representations in LQ settings.

Quantization-Aware Training (QAT) Hubara et al. (2018); Tailor et al. (2021); Zafrir et al. (2019); Chen et al. (2024); Mishchenko et al. (2019); Wei et al. (2025) integrates quantization into the training process to simulate low-bit computations during forward and backward passes He et al. (2024). Early QAT works Esser et al. (2020); Zhang et al. (2018) rely on the straight-through estimator (STE) Bengio et al. (2013) to approximate gradients of non-differentiable quantization functions, allowing end-to-end optimization. Recent works combined QAT with knowledge distillation Kim et al. (2019) and mixed-precision strategies Wang et al. (2019) to further reduce accuracy degradation, enabling robust deployment of convolutional and Transformer models under aggressive quantization constraints.

3 METHOD

(The quantization notation is described in Section A.) Our post-training procedure optimizes all quantization parameters exclusively via the relative entropy (KL divergence) between the full-precision (FP) and quan-

tized output distributions. Let $z_{FP}(x)$ and $z_{LQ}(x)$ denote the FP and quantized logits for input x . With temperature T and p as the probability of the model output, define

$$p_{FP}(x) := \text{softmax}(z_{FP}(x)/T), \quad p_{LQ}(x; \Theta) := \text{softmax}(z_{LQ}(x; \Theta)/T).$$

We determine the scale-factor for each tensor by minimizing the KL divergence on a small calibration set $\mathcal{X}$:

$$\mathcal{L}_{\text{KL.out}} = \frac{1}{|\mathcal{X}|} \sum_{x \in \mathcal{X}} \text{KL}(p_{FP}(x) \parallel p_{LQ}(x)) = \frac{1}{|\mathcal{X}|} \sum_{x \in \mathcal{X}} \sum_c p_{FP}^{(c)}(x) \log \frac{p_{FP}^{(c)}(x)}{p_{LQ}^{(c)}(x)}. \quad (1)$$

To avoid overfitting and preserve hardware-friendly ranges, we add a lightweight parameter regularizer:

$$\mathcal{L}_{\text{reg}} = \|\delta - \delta_0\|_2^2 + \|z - z_0\|_2^2, \quad (2)$$

where (δ_0, z_0) are the initial (calibration) estimates. The objective is to minimize $\mathcal{L}_{\text{CAT}}$ where,

$$\mathcal{L}_{\text{CAT}} = \mathcal{L}_{\text{KL.out}} + \lambda_p \mathcal{L}_{\text{reg}}. \quad (3)$$

This procedure ensures that each block is reconstructed to preserve the information content of its FP counterpart, reducing distributional mismatch layer by layer. The refined quantization parameters are then fixed for the remainder of the process, after which we apply logit-level transformation using CAT.

3.1 CLUSTER-BASED AFFINE TRANSFORMATION FOR ERROR REDUCTION

After refining quantization parameters, we address residual mismatches directly in the logit space through CAT (Figure 1). To make clustering more effective and computationally efficient, we use a *principal component analysis* (PCA) Wold et al. (1987) to decompose the LQ logits z_{LQ} . PCA reduces the dimensionality of the logits by discarding low-variance components, thereby (i) removing unnecessary complexity, (ii) improving cluster separability, and (iii) lowering clustering cost. Given calibration logits $\{z_{LQ}(x), z_{FP}(x)\}_{x \in \mathcal{X}}$, we apply PCA to reduce the dimensionality of $\{z_{LQ}(x)\}$ before clustering. Note that PCA is only used to obtain clusters. The affine transformation itself is always applied in the original logit space of dimension d . A clustering model is then fitted to the reduced features, yielding cluster assignments $c(x) \in \{1, \dots, K\}$. For each cluster $C_k = \{x : c(x) = k\}$, we seek Affine Parameters (AP_K) $(\gamma_k, \beta_k) \in \mathbb{R}^d$ (with d the dimensionality of logits) that best map quantized logits to their FP counterparts:

$$z_{FP}(x) \approx \gamma_k \odot z_{LQ}(x) + \beta_k, \quad x \in C_k. \quad (4)$$

Where $\odot$ is the element-wise (Hadamard) product. This can be derived in closed form by matching first- and second-order statistics of the two distributions:

$$\mu_{LQ,k} = \frac{1}{|C_k|} \sum_{x \in C_k} z_{LQ}(x), \quad \mu_{FP,k} = \frac{1}{|C_k|} \sum_{x \in C_k} z_{FP}(x). \quad (5)$$

Here, $\mu_{LQ,k}$ and $\mu_{FP,k}$ denote the mean logits of the LQ and FP models over cluster C_k , respectively.

$$\sigma_{LQ,k}^2 = \frac{1}{|C_k|} \sum_{x \in C_k} (z_{LQ}(x) - \mu_{LQ,k})^{\odot 2}. \quad (6)$$

$\sigma_{LQ,k}^2$ denotes the variance of LQ logits within cluster C_k .

$$\text{cov}_{LQ,FP,k} = \frac{1}{|C_k|} \sum_{x \in C_k} (z_{LQ}(x) - \mu_{LQ,k}) \odot (z_{FP}(x) - \mu_{FP,k}). \quad (7)$$

$\text{cov}_{LQ,FP,k}$ denotes the element-wise covariance between LQ and FP logits over cluster C_k . The element-wise affine parameters are then estimated as

$$\gamma_k = \frac{\text{cov}_{LQ,FP,k}}{\sigma_{LQ,k}^2 + \epsilon}, \quad \beta_k = \mu_{FP,k} - \gamma_k \odot \mu_{LQ,k}, \quad (8)$$

where division is elementwise and $\epsilon > 0$ avoids division by zero. Thus, for each cluster k , the transformation is obtained by aligning the cluster-wise mean and variance of quantized logits to those of FP, without requiring gradient optimization. We do not require backpropagation through the network to find (γ_k, β_k) . For each cluster k , we estimate γ_k and β_k as parameters to minimize the error between p_{LQ} and p_{FP} using Eq. (4). This black-box procedure fits the affine terms using only a small sample set, avoids gradient computations entirely, and naturally yields low-precision $\{\gamma_k\}_{k=1}^K$ and $\{\beta_k\}_{k=1}^K$ suitable for deployment. Our two-stage pipeline first optimizes equation 3 to refine quantization parameters using only the output-level KL loss. We then find γ and β through the CAT fitting phase. In inference, CAT requires only a single affine transformation per cluster assignment, introducing negligible overhead while substantially reducing the FP/LQ gap. We first assign the logits z_{LQ} to a cluster C_k . We then apply the following α -blended correction:

$$\tilde{z} = (1 - \alpha) z_{LQ} + \alpha (\gamma_k \odot z_{LQ} + \beta_k),$$

where $\alpha \in [0, 1]$ controls the contribution of the original quantized output and the affine transformed one. The pseudocode for CAT fitting and inference is detailed in Section B.

4 EXPERIMENTS

Setup. We use K-Means as the clustering algorithm. We evaluate CAT on ImageNet-1K Russakovsky et al. (2015) using ResNet-18/50 He et al. (2016), MobileNetV2 Sandler et al. (2018), RegNetX-600MF/3.2GF Radosavovic et al. (2020), and MNasX2 Tan et al. (2019) under different LQ settings, denoted as $\{\text{weight bit-width}\}A\{\text{activation bit-width}\}$: W4A4, W2A4, W4A2, and W2A2. CAT is compared against strong PTQ baselines (ACIQ-Mix, LAPQ, Bit-Split, AdaRound, QDrop, PD-Quant) with identical hyperparameters to ensure fairness. Models are quantized channel-wise, calibrated with AdaRound (20k iterations, batch size 64, 1,024 samples), and follow prior work by keeping the last layer at 8-bit. CAT adopts the same calibration as PD-Quant, with KL temperature 0.4 and learning rate 4×10^{-5} . All experiments are run on NVIDIA L40 GPUs, repeated with three seeds, and we report the mean and standard deviation of Top-1 accuracy. We ablate CAT hyperparameters (α , number of clusters (# Clusters), PCA dimension, and clustering samples) to analyze their effect on performance.

Quantitative Comparison against State-of-the-Art CAT provides advantages across all quantization settings and architectures (Table 2); however, the magnitude of improvements varies depending on the bit-width configuration and the network’s capacity. Overall, our results show that CAT provides the greatest benefits for networks with a larger gap between FP and LQ, such as MNasX2 (W2A2). In such settings, quantization errors accumulate heavily, and CAT’s cluster-specific affine correction substantially restores performance. With both weights and activations quantized to 2 bits (W2A2), CAT shows the best improvement. On ResNet-50, CAT achieves 58.08%, an improvement of +1.05% over PD-Quant. On MNasX2, CAT reaches 29.20%, outperforming PD-Quant by +1.25%. Gains are also consistent across ResNet-18 (+0.32%), MobileNetV2 (+0.51%), and RegNetX-3.2GF (+1.06%). Similarly, with 4-bit weights and 2-bit activations (W4A2), where the activation bottleneck is severe, CAT consistently delivers improvements. On ResNet-18, CAT yields 58.68%, exceeding PD-Quant by +0.11%, while on MNasX2, CAT improves accuracy to

Table 2: Top-1 accuracy (%) on ImageNet-1K for various PTQ methods across architectures.

Methods	Bits (W/A)	ResNet-18	ResNet-50	MobileNetV2	RegNetX-600MF	RegNetX-3.2GF	MNASX2
Full Prec.	32/32	71.01	76.63	72.62	73.52	78.46	76.52
ACIQ-Mix Banner et al. (2019)	4/4	67.00	73.80	—	—	—	—
LAPQ Nahshan et al. (2021)		60.30	70.00	49.70	57.71	55.89	65.32
Bit-Split Wang et al. (2020)		67.56	73.71	—	—	—	—
AdaRound Nagel et al. (2020)		67.96	73.88	61.52	68.20	73.85	68.86
QDrop Wei et al. (2022)		69.17	75.15	68.07	70.91	76.40	72.81
PD-Quant Liu et al. (2023)		69.14 $\pm$ 0.10	75.07 $\pm$ 0.09	68.18 $\pm$ 0.02	70.96 $\pm$ 0.03	76.54 $\pm$ 0.02	73.24 $\pm$ 0.02
Ours		69.18 $\pm$ 0.09	75.12 $\pm$ 0.07	68.22 $\pm$ 0.01	70.98 $\pm$ 0.01	76.61 $\pm$ 0.02	73.31 $\pm$ 0.05
LAPQ	2/4	0.18	0.14	0.13	0.17	0.12	0.18
Adaround		0.11	0.12	0.15	—	—	—
QDrop		64.57	70.09	53.37	63.18	71.96	63.23
PD-Quant		65.10 $\pm$ 0.02	70.84 $\pm$ 0.06	55.30 $\pm$ 0.24	63.92 $\pm$ 0.24	72.36 $\pm$ 0.13	63.32 $\pm$ 0.24
Ours		65.26 $\pm$ 0.06	71.29 $\pm$ 0.02	55.47 $\pm$ 0.22	64.2 $\pm$ 0.28	72.71 $\pm$ 0.12	63.96 $\pm$ 0.32
QDrop	4/2	57.56	63.26	17.30	49.73	62.79	34.12
PD-Quant		58.57 $\pm$ 0.18	64.24 $\pm$ 0.02	20.14 $\pm$ 0.52	51.17 $\pm$ 0.27	62.68 $\pm$ 0.08	39.43 $\pm$ 0.34
Ours		58.68 $\pm$ 0.15	64.38 $\pm$ 0.06	20.53 $\pm$ 0.53	51.52 $\pm$ 0.26	63.03 $\pm$ 0.05	40.14 $\pm$ 0.27
QDrop	2/2	51.42	55.45	10.28	39.01	54.38	23.59
PD-Quant		52.87 $\pm$ 0.03	57.03 $\pm$ 0.12	13.65 $\pm$ 0.63	40.71 $\pm$ 0.13	55.08 $\pm$ 0.13	27.95 $\pm$ 0.69
Ours		53.19 $\pm$ 0.07	58.08 $\pm$ 0.11	14.16 $\pm$ 0.61	41.38 $\pm$ 0.1	56.14 $\pm$ 0.09	29.20 $\pm$ 0.76

Note: PD-Quant results are reproduced from our runs.

40.14%, a substantial +0.71% increase. Gains are also observed on MobileNetV2 (+0.39%) and RegNetX-600MF (+0.35%). These findings demonstrate that W2A2 and W4A2 are the most error-prone regimes, and CAT’s targeted corrections are especially effective at resolving their distorted feature representations.

In another asymmetric setting, W2A4, CAT also provides consistent improvements. For instance, CAT improves over PD-Quant by +0.19% on ResNet-18, +0.45% on ResNet-50, and +0.35% on RegNetX-3.2GF. On MobileNetV2, CAT achieves 55.47%, improving upon PD-Quant by +0.17%, while MNasX2 benefits from a +0.64% gain. Interestingly, W2A4 tends to produce more diverse outcomes compared to W2A2 or W4A2: while low-bit weights distort the learned filters and increase the risk of incorrect feature extraction, the higher activation precision (4-bit) preserves a broader dynamic range, enabling richer but more variable behavior than the severely compressed 2-bit activation cases. At higher precision (W4A4), where the discrepancy between full-precision (FP) and low-bit (LQ) networks is relatively small, CAT provides only marginal improvements since its corrections are designed to bridge larger gaps. In this regime, CAT achieves accuracy on par with or slightly better than prior methods. On ResNet-18, CAT obtains 69.18%, essentially matching PD-Quant (69.14%), while on RegNetX-3.2GF, CAT improves to 76.61%, the best among all compared approaches. Similar trends are observed on MobileNetV2 (68.22%) and ResNet-50 (75.12%), confirming that when quantization noise is less severe, CAT closely mimics the FP network but cannot deliver substantial gains. An important observation in this is that networks with a larger discrepancy between FP and LQ outcomes benefit the most from CAT in LQ settings. This effect arises because low-capacity models have limited redundancy to absorb quantization errors, making CAT’s correction particularly impactful. Across all architectures and bit-widths, CAT either matches or outperforms prior state-of-the-art PTQ methods. The improvements are most pronounced under LQ settings (W2A2, W4A2, W2A4) and on compact models, where quantization errors are most destructive. CAT achieves more effective reduction than plain affine mapping, thereby establishing a new state-of-the-art in PTQ. Section I discusses a statistical analysis of CAT performance across the cross of model parameters and LQ settings.

Enhance PTQ methods with CAT Table 3 (More comprehensive results are provided in Appendix Table 5) provides a direct comparison of multiple PTQ baselines with and without the proposed CAT correction across different architectures and bit-width settings. The values in parentheses indicate the performance difference $\Delta = (\text{CAT} - \text{Base})$, where green arrows ($\uparrow$) mark improvements and red arrows ($\downarrow$) indicate degradations. Across nearly all architectures and methods, CAT consistently enhances performance. The

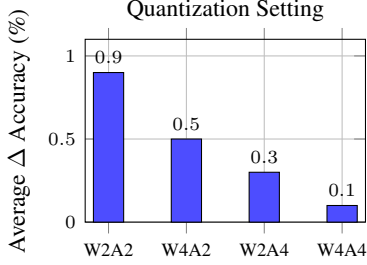


Figure 2: Average improvements $\Delta = (\text{CAT} - \text{Base})$ across quantization settings.

Table 3: ResNet-18 Top-1 Accuracy (%) with/without CAT. Δ shows improvement (green $\uparrow$) or degradation (red $\downarrow$).

	W4A4				W4A2			
	Adaround	BRECQ	QDrop	PD-Quant	Adaround	BRECQ	QDrop	PD-Quant
Base	2.59	69.32	69.17	69.14	1.18	49.44	57.91	58.57
+ CAT	2.58	69.35	69.24	69.18	1.14	50.49	58.12	58.68
Δ	-0.01 $\downarrow$	+0.03 $\uparrow$	+0.07 $\uparrow$	+0.04 $\uparrow$	-0.04 $\downarrow$	+1.05 $\uparrow$	+0.21 $\uparrow$	+0.11 $\uparrow$

	W2A4				W2A2			
	Adaround	BRECQ	QDrop	PD-Quant	Adaround	BRECQ	QDrop	PD-Quant
Base	18.48	62.69	64.52	65.10	2.59	40.78	51.47	52.86
+ CAT	22.04	62.67	64.76	65.35	2.58	41.75	51.76	53.19
Δ	+3.56 $\uparrow$	-0.02 $\downarrow$	+0.24 $\uparrow$	+0.25 $\uparrow$	-0.01 $\downarrow$	+0.97 $\uparrow$	+0.29 $\uparrow$	+0.32 $\uparrow$

improvements are particularly pronounced in LQ settings (W2A2 and W4A2), where quantization errors are most severe. For example, CAT boosts PD-Quant on ResNet-50 (W2A2) by +1.21%, on RegNetX-3.2GF (W2A2) by +1.01%, and on MNasX2 (W2A2) by +1.24%. Similarly, for W4A2, CAT improves QDrop on MNasX2 by more than +1% and PD-Quant on MobileNetV2 by +0.37%. These gains confirm that CAT is most effective in highly error-prone regimes, where its cluster-aware affine correction recovers distorted representations. In asymmetric quantization (W2A4), CAT again provides steady improvements, albeit with smaller margins (e.g., ResNet-50 +0.42%, RegNetX-3.2GF +0.45%), while in higher-precision settings (W4A4), the effect is marginal but consistently non-negative, reflecting that the baseline already closely approximates full-precision. An important observation is that smaller-capacity architectures exhibit disproportionately large gains from CAT, providing new evidence that networks with a large FP-LQ gap benefit the most in LQ settings. For instance, LQ MNasX2 gains more than +1%, demonstrating that CAT is particularly valuable for compact models that lack redundancy to absorb quantization errors. Overall, Table 3 highlights CAT’s robustness as a drop-in enhancement: it supplements diverse PTQ baselines (Adaround, BRECQ, Qdrop, PD-Quant) consistently yielding improved accuracy while never severely degrading performance. This consistency demonstrates CAT’s generality and compatibility with existing PTQ pipelines. Figure 2 shows the average improvement of CAT over the Base method. The largest gain is observed for W2A2 quantization (0.9%) across all PTQ methods and models. The second-highest improvement occurs for W4A2 (0.5%), highlighting the effectiveness of CAT under very low-bit activation quantization. For W2A4 and W4A4, the accuracy increases by 0.3% and 0.1%, respectively, demonstrating the generalization capability of CAT. Notably, 2-bit activation settings benefit the most from CAT, since their severely limited precision reduces the ability to preserve information entropy, making them more reliant on CAT’s error reduction mechanism. Section H represents LQ ViT error reduction using CAT.

5 ABLATION STUDY

Ablation on Blending Coefficient α . This study ablates the effect of the blending coefficient α , where $\alpha = 0$ corresponds to using only PTQ logits without any CAT correction, and $\alpha = 1$ corresponds to relying entirely on CAT-corrected logits without involving the original PTQ. As shown in Fig. 3, varying α directly influences the final accuracy across different bit-width regimes (The comprehensive ablation of Blending Coefficient α for different LQ settings and architectures provided in Section D).

For the 2-bit activation quantization (W2A2 and W4A2), moderate blending ($\alpha \approx 0.3$ – 0.4) consistently provides the highest performance, indicating that CAT is most effective when used as a supplement to the baseline quantized outputs rather than a full replacement. In the asymmetric regime (W2A4), accuracy is relatively stable across a broad range of α , reflecting that higher-precision activations reduce sensitivity to

blending, although small gains are still observed around $\alpha \approx 0.4$. At higher precision (W4A4), the performance is nearly flat across all values of α , as the quantized model is already close to full-precision accuracy. Overall, these results demonstrate that CAT reliably improves PTQ performance, especially in LQ settings, but that using CAT alone ($\alpha = 1$) is suboptimal.

Ablation on the Number of Clusters. We next ablate the influence of the number of clusters used in CAT, as shown in Fig. 4 for ResNet-18 W2A4 with PCA dimension fixed to 50 and blending coefficient $\alpha = 0.6$ (The comprehensive ablation of the number of clusters for different LQ settings and architectures is provided in Section E). For W2A2, performance is highest with a small number of clusters and gradually declines as the cluster count increases. This suggests that in extremely quantized regimes, a compact cluster partition provides stable corrections, while too many clusters lead to overfitting of noise in the heavily distorted logit space. For W4A2, a similar trend is observed: accuracy peaks at very low cluster counts (1–8 clusters) and then decreases slowly with more clusters, indicating diminishing returns once the coarse logit structure has been captured. In the asymmetric case W2A4, accuracy remains largely stable across a broad range of cluster counts, confirming that higher-precision activations mitigate sensitivity to clustering granularity. At higher precision (W4A4), accuracy is nearly unaffected by number of clusters, as the quantized logits already closely approximate the full-precision distribution. Overall, these results demonstrate that the optimal number of clusters is inherently fitted to the nature of quantization precision: lower precision reduces diversity in the logit space and thus favors fewer clusters, whereas higher precision allows richer structures that can benefit from larger cluster counts. This further suggests that clustering is not only a useful but also an essential component of the affine transformation, enabling it to adaptively reduce quantization errors according to the underlying representation capacity of the quantized network.

Ablation on PCA Dimension k . We further study the effect of the PCA dimension k used to reduce the logit space before clustering, as shown in Fig. 5 for ResNet-18 under the W2A4 quantization setting (The comprehensive ablation of the PCA dimension for different LQ settings and architectures is provided in Section F).

For W2A2, performance is maximized when k is very small (1–5) and gradually decreases as k increases. This indicates that in extremely quantized networks, only the coarse structure of the logits can be reliably captured, and projecting onto a compact subspace avoids fitting noise.

For W4A2, a similar but weaker trend is observed: accuracy peaks when k is kept small (≤ 10) and remains stable for moderate values before slightly degrading with very large k . In the asymmetric W2A4 setting, accuracy is largely stable across the full range of k , suggesting that higher activation precision preserves sufficient feature variability to tolerate richer subspaces without significant overfitting.

At higher precision (W4A4), accuracy is almost invariant to the PCA dimension, as the quantized logits already closely match the full-precision distribution and PCA reduction plays only a minor role.

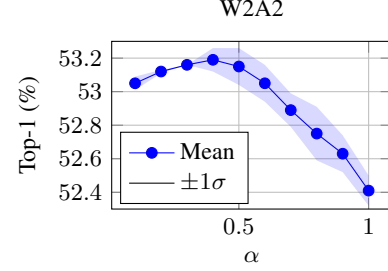


Figure 3: Top-1 accuracy for W2A2 across different α values with mean and $\pm 1\sigma$ range.

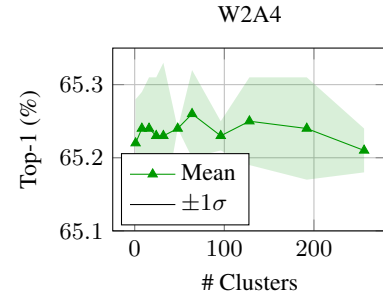


Figure 4: Top-1 accuracy for W2A4 across different number of Clusters (mean $\pm 1\sigma$).

Results indicate that the optimal PCA dimension depends on the severity of quantization: lower-precision networks benefit from aggressive dimensionality reduction that filters out noise, while higher-precision networks can afford larger k without degradation. This highlights PCA as an essential regularization step that adapts the clustering space to the effective diversity of the quantized representations.

Ablation on CAT fitting sample size. Table 4 studies the impact of the number of samples for the fitting of CAT on ResNet-18 on the W2A2 setting (The comprehensive ablation of sample size on CAT’s performance for different ResNet-18 LQ settings is provided in Section G). We observe a consistent monotonic trend with diminishing returns: accuracy improves rapidly when increasing samples from 10 to 500–1000, and then plateaus. Under W2A2, the mean Top-1 rises from 46.48% (10 samples) to 53.04% (1,000), a gain of $\approx +6.6$ points, with only marginal changes beyond 1,000 (e.g., 53.15% at 100,000). W4A2 shows a similar but smaller effect (53.67% $\rightarrow$ 58.54%, $\approx +4.9$). In contrast, higher-precision settings are less sensitive: W2A4 improves by $\approx +2.7$ (62.60% $\rightarrow$ 65.33%), while W4A4 gains only $\approx +1.1$ (68.13% $\rightarrow$ 69.26%). The shaded $\pm 1\sigma$ bands narrow as sample size grows, indicating more stable calibration and reduced run-to-run variability, especially pronounced in W2A2/W4A2. A calibration set of ~ 500 –1,000 samples captures nearly all attainable gains across regimes, striking a favorable accuracy/cost trade-off. Using $\geq 10k$ samples yields negligible additional improvements, particularly in W4A4, where performance is near its ceiling.

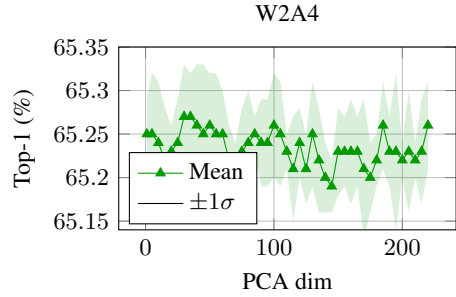


Figure 5: Top-1 accuracy for W2A4 across different PCA dim (mean $\pm 1\sigma$).

6 LIMITATIONS

Table 4: Top-1 accuracy for W2A2 across different numbers of samples CAT fitting.

# Samples	Top-1 (%)
10	46.48
50	51.79
100	52.49
500	53.02
1K	53.04
10K	53.15
100K	53.14

While CAT consistently improves accuracy over PTQ baselines, it introduces additional parameters due to the clustering step and the cluster-specific affine corrections. In particular, PTQ baselines do not maintain any auxiliary parameters beyond the quantized model itself, whereas CAT requires storing the clustering model and the affine coefficients (γ, β) for each cluster. Here, we analyse additional parameters by considering the K-Means clustering algorithm with k means. This overhead grows linearly with the number of clusters k and the logit dimensionality d (equal to the number of classes). The additional parameter count of CAT can therefore be approximated as $\text{CAT}_{\text{\#params}} = (k \times d \text{ for cluster-wise affine coefficients}) + (k \times d \text{ for } k\text{-means centroids})$. For example, for ResNet-18 with 11.6 million parameters, the number of model parameters in addition to CAT with $k = 50$ and $d = 1000$ is $11,600,000 + 2 * (50 * 1000) = 11,700,000$. In this example, CAT adds only $\sim 0.9\%$ parameter overhead relative to the baseline model. Nevertheless, this extra storage may be undesirable for extremely resource-constrained deployments, which we identify as a limitation of CAT compared to PTQ baselines.

7 CONCLUSION

We studied the gap between full-precision (FP) and low-bit quantized (LQ) networks, focusing on reducing quantization errors through affine transformations. Our initial findings showed that a plain affine map-

ping with a uniform parameter set can even worsen PTQ performance. To overcome this, we introduced Cluster-based Affine Transformation (CAT), which leverages the clusterability of quantized logits and applies cluster-specific affine corrections. CAT consistently improves accuracy under low-bit settings with only negligible parameter overhead, achieving state-of-the-art performance on ImageNet-1K. These results highlight that cluster-aware error reduction is a powerful and generalizable strategy for enhancing PTQ, particularly in extremely low-bit regimes.

REFERENCES

- Ron Banner, Yury Nahshan, and Daniel Soudry. Post training 4-bit quantization of convolutional networks for rapid-deployment. In H. Wallach, H. Larochelle, A. Beygelzimer, F. d'Alché-Buc, E. Fox, and R. Garnett (eds.), *Advances in Neural Information Processing Systems*, volume 32. Curran Associates, Inc., 2019. URL https://proceedings.neurips.cc/paper_files/paper/2019/file/c0a62e133894cdce435bcb4a5df1db2d-Paper.pdf.
- Yoshua Bengio, Nicholas Léonard, and Aaron Courville. Estimating or propagating gradients through stochastic neurons for conditional computation. *arXiv preprint arXiv:1308.3432*, 2013.
- Yaohui Cai, Zhewei Yao, Zhen Dong, Amir Gholami, Michael W. Mahoney, and Kurt Keutzer. Zeroq: A novel zero shot quantization framework. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, June 2020.
- Lei Chen, Yuan Meng, Chen Tang, Xinzhu Ma, Jingyan Jiang, Xin Wang, Zhi Wang, and Wenwu Zhu. Q-dit: Accurate post-training quantization for diffusion transformers. In *Proceedings of the Computer Vision and Pattern Recognition Conference (CVPR)*, pp. 28306–28315, June 2025.
- Mengzhao Chen, Wenqi Shao, Peng Xu, Jiahao Wang, Peng Gao, Kaipeng Zhang, Yu Qiao, and Ping Luo. Efficientqat: Efficient quantization-aware training for large language models. *CoRR*, abs/2407.11062, 2024. URL <https://doi.org/10.48550/arXiv.2407.11062>.
- Tianshu Chu, Zuopeng Yang, and Xiaolin Huang. Improving the post-training neural network quantization by prepositive feature quantization. *IEEE Transactions on Circuits and Systems for Video Technology*, 34(4):3056–3060, 2024. doi: 10.1109/TCSVT.2023.3311923.
- Xin Ding, Xiaoyu Liu, Zhijun Tu, Yun Zhang, Wei Li, Jie Hu, Hanting Chen, Yehui Tang, Zhiwei Xiong, Baoqun Yin, and Yunhe Wang. CBQ: Cross-block quantization for large language models. In *The Thirteenth International Conference on Learning Representations*, 2025. URL <https://openreview.net/forum?id=eW4yh6HKz4>.
- Yifu Ding, Haotong Qin, Qinghua Yan, Zhenhua Chai, Junjie Liu, Xiaolin Wei, and Xianglong Liu. Towards accurate post-training quantization for vision transformer. In *Proceedings of the 30th ACM International Conference on Multimedia*, MM '22, pp. 5380–5388, New York, NY, USA, 2022. Association for Computing Machinery. ISBN 9781450392037. doi: 10.1145/3503161.3547826. URL <https://doi.org/10.1145/3503161.3547826>.
- Steven K. Esser, Jeffrey L. McKinstry, Deepika Bablani, Rathinakumar Appuswamy, and Dharmendra S. Modha. Learned step size quantization. In *International Conference on Learning Representations*, 2020. URL <https://openreview.net/forum?id=rkgO66VKDS>.
- Ruihao Gong, Xianglong Liu, Yuhang Li, Yunqiang Fan, Xiuying Wei, and Jinyang Guo. Pushing the limit of post-training quantization. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 47(7):5556–5570, 2025. doi: 10.1109/TPAMI.2025.3554523.
- Simla Burcu Harma, Ayan Chakraborty, Elizaveta Kostenok, Danila Mishin, Dongho Ha, Babak Falsafi, Martin Jaggi, Ming Liu, Yunho Oh, Suvinay Subramanian, and Amir Yazdanbakhsh. Effective interplay between sparsity and quantization: From theory to practice. In *The Thirteenth International Conference on Learning Representations*, 2025. URL <https://openreview.net/forum?id=wJv4AIt4sK>.
- Kaiming He, Xiangyu Zhang, Shaoqing Ren, and Jian Sun. Deep residual learning for image recognition. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, June 2016.

- Yefei He, Jing Liu, Weijia Wu, Hong Zhou, and Bohan Zhuang. EfficientDM: Efficient quantization-aware fine-tuning of low-bit diffusion models. In *The Twelfth International Conference on Learning Representations*, 2024. URL <https://openreview.net/forum?id=UmMa3UNDaz>.
- Itay Hubara, Matthieu Courbariaux, Daniel Soudry, Ran El-Yaniv, and Yoshua Bengio. Quantized neural networks: Training neural networks with low precision weights and activations. *Journal of Machine Learning Research*, 18(187):1–30, 2018. URL <http://jmlr.org/papers/v18/16-456.html>.
- Yongkweon Jeon, Chungman Lee, Eulrang Cho, and Yeonju Ro. Mr.biq: Post-training non-uniform quantization based on minimizing the reconstruction error. In *2022 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pp. 12319–12328, 2022. doi: 10.1109/CVPR52688.2022.01201.
- Jangho Kim, Yash Bhargat, Jinwon Lee, Chirag Patel, and Nojun Kwak. Qkd: Quantization-aware knowledge distillation. *arXiv preprint arXiv:1911.12491*, 2019.
- Jemin Lee, Yongin Kwon, Sihyeong Park, Misun Yu, Jeman Park, and Hwanjun Song. Q-hyvit: Post-training quantization of hybrid vision transformers with bridge block reconstruction for iot systems. *IEEE Internet of Things Journal*, 11(22):36384–36396, 2024. doi: 10.1109/IJOT.2024.3403844.
- Muyang Li, Yujun Lin, Zhekai Zhang, Tianle Cai, Junxian Guo, Xiuyu Li, Enze Xie, Chenlin Meng, Jun-Yan Zhu, and Song Han. SVDQuant: Absorbing outliers by low-rank component for 4-bit diffusion models. In *The Thirteenth International Conference on Learning Representations*, 2025a. URL <https://openreview.net/forum?id=vWR3KuiQur>.
- Yixiao Li, Yifan Yu, Chen Liang, Nikos Karampatziakis, Pengcheng He, Weizhu Chen, and Tuo Zhao. Loftq: LoRA-fine-tuning-aware quantization for large language models. In *The Twelfth International Conference on Learning Representations*, 2024. URL <https://openreview.net/forum?id=LzPWWPAdY4>.
- Yuhang Li, Ruihao Gong, Xu Tan, Yang Yang, Peng Hu, Qi Zhang, Fengwei Yu, Wei Wang, and Shi Gu. {BRECQ}: Pushing the limit of post-training quantization by block reconstruction. In *International Conference on Learning Representations*, 2021. URL <https://openreview.net/forum?id=POWv6hDd9XH>.
- Zhikai Li, Junrui Xiao, Lianwei Yang, and Qingyi Gu. Repq-vit: Scale reparameterization for post-training quantization of vision transformers. In *Proceedings of the IEEE/CVF International Conference on Computer Vision (ICCV)*, pp. 17227–17236, October 2023.
- Ziqi Li, Tao Gao, Yisheng An, Ting Chen, Jing Zhang, Yuanbo Wen, Mengkun Liu, and Qianxi Zhang. Brain-inspired spiking neural networks for energy-efficient object detection. In *Proceedings of the Computer Vision and Pattern Recognition Conference (CVPR)*, pp. 3552–3562, June 2025b.
- Yang Lin, Tianyu Zhang, Peiqin Sun, Zheng Li, and Shuchang Zhou. Fq-vit: Post-training quantization for fully quantized vision transformer. *arXiv preprint arXiv:2111.13824*, 2021.
- Jiawei Liu, Lin Niu, Zhihang Yuan, Dawei Yang, Xinggang Wang, and Wenyu Liu. Pd-quant: Post-training quantization based on prediction difference metric. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pp. 24427–24437, June 2023.
- Yuxiao Ma, Huixia Li, Xiawu Zheng, Feng Ling, Xuefeng Xiao, Rui Wang, Shilei Wen, Fei Chao, and Rongrong Ji. Affinequant: Affine transformation quantization for large language models. In *The Twelfth International Conference on Learning Representations*, 2024. URL <https://openreview.net/forum?id=of2rhALq8l>.

- Yuriy Mishchenko, Yusuf Goren, Ming Sun, Chris Beauchene, Spyros Matsoukas, Oleg Rybakov, and Shiv Naga Prasad Vitaladevuni. Low-bit quantization and quantization-aware training for small-footprint keyword spotting. In *2019 18th IEEE International Conference On Machine Learning And Applications (ICMLA)*, pp. 706–711, 2019. doi: 10.1109/ICMLA.2019.00127.
- Markus Nagel, Rana Ali Amjad, Mart Van Baalen, Christos Louizos, and Tijmen Blankevoort. Up or down? Adaptive rounding for post-training quantization. In Hal Daumé III and Aarti Singh (eds.), *Proceedings of the 37th International Conference on Machine Learning*, volume 119 of *Proceedings of Machine Learning Research*, pp. 7197–7206. PMLR, 13–18 Jul 2020. URL <https://proceedings.mlr.press/v119/nagel20a.html>.
- Yury Nahshan, Brian Chmiel, Chaim Baskin, Evgenii Zheltonozhskii, Ron Banner, Alex M Bronstein, and Avi Mendelson. Loss aware post-training quantization. *Machine Learning*, 110(11):3245–3262, 2021.
- Ting Qin, Zhao Li, Jiaqi Zhao, Yuting Yan, and Yafei Du. Mixed precision quantization based on information entropy. *Scientific Reports*, 15(1):12974, 2025.
- Ilija Radosavovic, Raj Prateek Kosaraju, Ross Girshick, Kaiming He, and Piotr Dollar. Designing network design spaces. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, June 2020.
- Nikhila Ravi, Valentin Gabeur, Yuan-Ting Hu, Ronghang Hu, Chaitanya Ryali, Tengyu Ma, Haitham Khedr, Roman Rädle, Chloe Rolland, Laura Gustafson, Eric Mintun, Junting Pan, Kalyan Vasudev Alwala, Nicolas Carion, Chao-Yuan Wu, Ross Girshick, Piotr Dollar, and Christoph Feichtenhofer. SAM 2: Segment anything in images and videos. In *The Thirteenth International Conference on Learning Representations*, 2025. URL <https://openreview.net/forum?id=Ha6RTeWmd0>.
- Olga Russakovsky, Jia Deng, Hao Su, Jonathan Krause, Sanjeev Satheesh, Sean Ma, Zhiheng Huang, Andrej Karpathy, Aditya Khosla, Michael Bernstein, Alexander C. Berg, and Li Fei-Fei. ImageNet Large Scale Visual Recognition Challenge. *International Journal of Computer Vision (IJCV)*, 115(3):211–252, 2015. doi: 10.1007/s11263-015-0816-y.
- Mark Sandler, Andrew Howard, Menglong Zhu, Andrey Zhmoginov, and Liang-Chieh Chen. Mobilenetv2: Inverted residuals and linear bottlenecks. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, June 2018.
- Utkarsh Saxena, Sayeh Sharify, Kaushik Roy, and Xin Wang. Resq: Mixed-precision quantization of large language models with low-rank residuals. In *Forty-second International Conference on Machine Learning*, 2025. URL <https://openreview.net/forum?id=4qIP1sXcR1>.
- Xuan Shen, Weize Ma, Jing Liu, Changdi Yang, Rui Ding, Quanyi Wang, Henghui Ding, Wei Niu, Yanzhi Wang, Pu Zhao, Jun Lin, and Jiuxiang Gu. Quartdepth: Post-training quantization for real-time depth estimation on the edge. In *Proceedings of the Computer Vision and Pattern Recognition Conference (CVPR)*, pp. 11448–11460, June 2025.
- Junqi Shi, Zhujia Chen, Hanfei Li, Qi Zhao, Ming Lu, Tong Chen, and Zhan Ma. On quantizing neural representation for variable-rate video coding. In *The Thirteenth International Conference on Learning Representations*, 2025. URL <https://openreview.net/forum?id=44cMlQSreK>.
- Zhenhong Sun, Ce Ge, Junyan Wang, Ming Lin, Hesen Chen, Hao Li, and Xiuyu Sun. Entropy-driven mixed-precision quantization for deep network design. In S. Koyejo, S. Mohamed, A. Agarwal, D. Belgrave, K. Cho, and A. Oh (eds.), *Advances in Neural Information Processing Systems*, volume 35, pp. 21508–21520. Curran Associates, Inc., 2022. URL https://proceedings.neurips.cc/paper_files/paper/2022/file/86e7eb16d33d59e62d1b0a079ea058d-Paper-Conference.pdf.

- Shyam Anil Tailor, Javier Fernandez-Marques, and Nicholas Donald Lane. Degree-quant: Quantization-aware training for graph neural networks. In *International Conference on Learning Representations*, 2021. URL <https://openreview.net/forum?id=NSBrFgJAHg>.
- Mingxing Tan, Bo Chen, Ruoming Pang, Vijay Vasudevan, Mark Sandler, Andrew Howard, and Quoc V. Le. Mnasnet: Platform-aware neural architecture search for mobile. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, June 2019.
- Kuan Wang, Zhijian Liu, Yujun Lin, Ji Lin, and Song Han. Haq: Hardware-aware automated quantization with mixed precision. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, June 2019.
- Peisong Wang, Qiang Chen, Xiangyu He, and Jian Cheng. Towards accurate post-training network quantization via bit-split and stitching. In Hal Daumé III and Aarti Singh (eds.), *Proceedings of the 37th International Conference on Machine Learning*, volume 119 of *Proceedings of Machine Learning Research*, pp. 9847–9856. PMLR, 13–18 Jul 2020. URL <https://proceedings.mlr.press/v119/wang20c.html>.
- Wenbin Wang, Yongcheng Jing, Liang Ding, Yingjie Wang, Li Shen, Yong Luo, Bo Du, and Dacheng Tao. Retrieval-augmented perception: High-resolution image perception meets visual RAG. In *Forty-second International Conference on Machine Learning*, 2025. URL <https://openreview.net/forum?id=X9vBykZVYg>.
- Quan Wei, Chung-Yiu Yau, Hoi To Wai, Yang Zhao, Dongyeop Kang, Youngsuk Park, and Mingyi Hong. RoSTE: An efficient quantization-aware supervised fine-tuning approach for large language models. In *Forty-second International Conference on Machine Learning*, 2025. URL <https://openreview.net/forum?id=h30EzoI3s0>.
- Xiuying Wei, Ruihao Gong, Yuhang Li, Xianglong Liu, and Fengwei Yu. QDrop: Randomly dropping quantization for extremely low-bit post-training quantization. In *International Conference on Learning Representations*, 2022. URL <https://openreview.net/forum?id=ySQH0oDyp7>.
- Eric W Weisstein. Affine transformation. <https://mathworld.wolfram.com/>, 2004.
- Svante Wold, Kim Esbensen, and Paul Geladi. Principal component analysis. *Chemometrics and Intelligent Laboratory Systems*, 2(1):37–52, 1987. ISSN 0169-7439. doi: [https://doi.org/10.1016/0169-7439\(87\)80084-9](https://doi.org/10.1016/0169-7439(87)80084-9). URL <https://www.sciencedirect.com/science/article/pii/0169743987800849>. Proceedings of the Multivariate Statistical Workshop for Geologists and Geochemists.
- Zhuguanyu Wu, Jiayi Zhang, Jiaxin Chen, Jinyang Guo, Di Huang, and Yunhong Wang. Aphq-vit: Post-training quantization with average perturbation hessian based reconstruction for vision transformers. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pp. 9686–9695, June 2025.
- Bohan Xiao, Peiyong Wang, Qisheng He, and Ming Dong. Deterministic image-to-image translation via denoising brownian bridge models with dual approximators. In *Proceedings of the Computer Vision and Pattern Recognition Conference (CVPR)*, pp. 28232–28241, June 2025.
- Zhihang Yuan, Chenhao Xue, Yiqi Chen, Qiang Wu, and Guangyu Sun. Ptq4vit: Post-training quantization for vision transformers with twin uniform quantization. In Shai Avidan, Gabriel Brostow, Moustapha Cissé, Giovanni Maria Farinella, and Tal Hassner (eds.), *Computer Vision – ECCV 2022*, pp. 191–207, Cham, 2022. Springer Nature Switzerland. ISBN 978-3-031-19775-8.

- Ofir Zafrir, Guy Boudoukh, Peter Izsak, and Moshe Wasserblat. Q8bert: Quantized 8bit bert. In *2019 Fifth Workshop on Energy Efficient Machine Learning and Cognitive Computing - NeurIPS Edition (EMC2-NIPS)*, pp. 36–39, 2019. doi: 10.1109/EMC2-NIPS53020.2019.00016.
- Maxime Zanella, Benoît Gérin, and Ismail Ben Ayed. Boosting vision-language models with transduction. In *The Thirty-eighth Annual Conference on Neural Information Processing Systems*, 2024. URL <https://openreview.net/forum?id=go4zzXBWVs>.
- Dongqing Zhang, Jiaolong Yang, Dongqiangzi Ye, and Gang Hua. Lq-nets: Learned quantization for highly accurate and compact deep neural networks. In *Proceedings of the European Conference on Computer Vision (ECCV)*, September 2018.
- Yunshan Zhong, You Huang, Jiawei Hu, Yuxin Zhang, and Rongrong Ji. Towards accurate post-training quantization of vision transformers via error reduction. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 47(4):2676–2692, 2025. doi: 10.1109/TPAMI.2025.3528042.
- Sifan Zhou, Zhihang Yuan, Dawei Yang, Xing Hu, Jian Qian, and Ziyu Zhao. Pillarhist: A quantization-aware pillar feature encoder based on height-aware histogram. In *Proceedings of the Computer Vision and Pattern Recognition Conference (CVPR)*, pp. 27336–27345, June 2025.
- Libo Zhu, Jianze Li, Haotong Qin, Wenbo Li, Yulun Zhang, Yong Guo, and Xiaokang Yang. Passionsr: Post-training quantization with adaptive scale in one-step diffusion based image super-resolution. In *Proceedings of the Computer Vision and Pattern Recognition Conference (CVPR)*, pp. 12778–12788, June 2025.

APPENDIX FOR CAT

A QUANTIZATION NOTATION

To better ground the discussion of CAT, we start by defining the notation used in neural network quantization. Quantization is the process of compacting a neural network by reducing the precision of full-precision weights and activations, considering the given bit-width allowance. For notational simplicity, we consider the quantization of a scalar value. Let $f \in \mathbb{R}$ denote a floating-point value (weight or activation) with real range

$$f_{\min} = \min(f), \quad f_{\max} = \max(f).$$

For a given bit-width b , the representable integer range is defined as

$$q_{\min} = -2^{b-1}, \quad q_{\max} = 2^{b-1} - 1,$$

where $q_{\min}$ and $q_{\max}$ denote the minimum and maximum values of the quantized representation. Assuming *scale factor* δ and *zero-point* z as parameters to convert floating-point value f to its quantized value q with N -bit representation, then we have,

$$\delta = \frac{f_{\max} - f_{\min}}{q_{\max} - q_{\min}}, \quad z = \left\lfloor q_{\min} - \frac{f_{\min}}{\delta} \right\rfloor.$$

The quantization and dequantization mappings are then given by

$$q = \text{clip}\left(\left\lfloor \frac{f}{\delta} \right\rfloor + z, q_{\min}, q_{\max}\right),$$

$$\hat{f} = \delta \cdot (q - z),$$

$\hat{f}$ is the reconstructed quantized value, and $\text{clip}(\cdot)$ ensures $q \in [q_{\min}, q_{\max}]$. Weights quantization is simpler since their values remain constant during PTQ, whereas activations quantization requires estimating their minimum and maximum. In post-training quantization (PTQ), $f_{\min}$ and $f_{\max}$ are estimated from a small calibration dataset. Different strategies include *min-max calibration*, which directly uses observed extrema, *percentile or clipping methods*, which discard outliers, and *optimization-based approaches* such as AdaRound or BRECQ, which refine δ (and sometimes z) by minimizing reconstruction error between FP and quantized outputs. We denote a quantization configuration by $WbAb$, where Wb and Ab specify the bit-widths used for weights and activations, respectively (e.g., $W2A4$ indicates 2-bit weights with 4-bit activations).

B CLUSTER-BASED AFFINE TRANSFORMATION (CAT) ALGORITHM

Algorithm 1 CAT Fitting

Require: $\text{all_q} \in \mathbb{R}^{N \times C}$, $\text{all_fp} \in \mathbb{R}^{N \times C}$, Clusters , pca_dim
Ensure: CLUSTER-MODEL, PCA, $\{\gamma_c\}_{c=1}^K$, $\{\beta_c\}_{c=1}^K$

- 1: $q_features \leftarrow \text{PCA.decomposition}(\text{pca_dim}, \text{all_q})$
- 2: $\text{cluster_ids} \leftarrow \text{CLUSTER-MODEL.fit}(\text{Clusters}, q_features)$
- 3: **for** $cid \leftarrow 1 \rightarrow \text{Clusters}$ **do**
- 4: $\text{idxs} \leftarrow \{i \mid \text{cluster_ids}[i] = cid\}$
- 5: **if** $|\text{idxs}| = 0$ **then**
- 6: $\gamma_{cid} \leftarrow \mathbf{1}_C$; $\beta_{cid} \leftarrow \mathbf{0}_C$
- 7: **continue**
- 8: **end if**
- 9: $Q \leftarrow \text{all_q}[\text{idxs}]$, $F \leftarrow \text{all_fp}[\text{idxs}]$
- 10: $\mu_Q \leftarrow \text{Mean}(Q)$, $\mu_F \leftarrow \text{Mean}(F)$
- 11: $\sigma_Q^2 \leftarrow \max(\text{Var}(Q), 10^{-8})$
- 12: $\gamma_{cid} \leftarrow \frac{\text{Mean}[(Q - \mu_Q) \odot (F - \mu_F)]}{\sigma_Q^2}$
- 13: $\beta_{cid} \leftarrow \mu_F - \gamma_{cid} \odot \mu_Q$
- 14: **end for**
- 15: **return** CLUSTER-MODEL, $\{\gamma_c\}$, $\{\beta_c\}$

Algorithm 2 CAT Inferencing

Require: q_logits , CLUSTER-MODEL, $\{\gamma_c\}$, $\{\beta_c\}$, α
Ensure: corrected logits, PCA

- 1: $q_feature \leftarrow \text{PCA.decomposition}(q_np)$
- 2: $\text{cluster_ids} \leftarrow \text{CLUSTER-MODEL.predict}(q_feature)$
- 3: $\text{corrected} \leftarrow []$
- 4: **for** $i = 1 \rightarrow |q_logits|$ **do**
- 5: $q \leftarrow q_logits[i]$
- 6: $cid \leftarrow \text{cluster_ids}[i]$
- 7: $\gamma \leftarrow \gamma_{cid}$, $\beta \leftarrow \beta_{cid}$
- 8: $\text{affine_corrected} \leftarrow q \cdot \gamma + \beta$
- 9: $\text{blended} \leftarrow q + \alpha \cdot (\text{affine_corrected} - q)$
- 10: Append blended to corrected
- 11: **end for**
- 12: **return** corrected

C ENHANCE PTQ METHODS WITH CAT

Table 5: Comparison of Quantization Methods (with/without CAT) across Architectures on ImageNet-1K. Values in parentheses show $\Delta = (\text{CAT} - \text{Base})$. **Green $\uparrow$** : CAT improves. **Red $\downarrow$** : CAT worse. The comprehensive description of this table’s results is provided in ??

Bits (W/A)	Methods	ResNet-18	ResNet-50	MobileNetV2	RegNetX-600MF	RegNetX-3.2GF	MNasX2
4/4	Adaround + CAT (Δ)	2.59 2.58 (-0.01 $\downarrow$)	1.57 1.57 (0.00)	24.22 28.51 (+4.29 $\uparrow$)	–	0.95 0.98 (+0.03 $\uparrow$)	0.68 0.75 (+0.07 $\uparrow$)
4/4	BRECQ + CAT (Δ)	69.32 69.35 (+0.03 $\uparrow$)	74.45 74.48 (+0.03 $\uparrow$)	4.56 6.48 (+1.92 $\uparrow$)	69.67 69.72 (+0.05 $\uparrow$)	74.57 74.68 (+0.11 $\uparrow$)	66.26 66.57 (+0.31 $\uparrow$)
4/4	Qdrop + CAT (Δ)	69.17 69.24 (+0.07 $\uparrow$)	75.08 75.18 (+0.10 $\uparrow$)	67.86 68.01 (+0.15 $\uparrow$)	70.78 70.89 (+0.11 $\uparrow$)	76.31 76.46 (+0.15 $\uparrow$)	72.91 73.04 (+0.13 $\uparrow$)
4/4	PD-Quant + CAT (Δ)	69.17 69.27 (+0.10 $\uparrow$)	–	68.20 68.28 (+0.08 $\uparrow$)	70.9 71.08 (+0.18 $\uparrow$)	–	–
4/2	Adaround + CAT (Δ)	1.18 1.14 (-0.04 $\downarrow$)	–	24.22 28.51 (+4.29 $\uparrow$)	–	–	0.15 0.17 (+0.02 $\uparrow$)
4/2	BRECQ + CAT (Δ)	49.44 50.49 (+1.05 $\uparrow$)	36.59 38.47 (+1.88 $\uparrow$)	0.23 0.43 (+0.20 $\uparrow$)	6.79 11.22 (+4.43 $\uparrow$)	6.48 10.16 (+3.68 $\uparrow$)	0.82 1.73 (+0.91 $\uparrow$)
4/2	Qdrop + CAT (Δ)	57.91 58.12 (+0.21 $\uparrow$)	63.15 63.54 (+0.39 $\uparrow$)	16.56 16.93 (+0.37 $\uparrow$)	49.83 50.05 (+0.22 $\uparrow$)	62.15 62.70 (+0.55 $\uparrow$)	30.43 31.44 (+1.01 $\uparrow$)
4/2	PD-Quant + CAT (Δ)	58.66 58.80 (+0.14 $\uparrow$)	64.26 64.49 (+0.23 $\uparrow$)	19.84 20.21 (+0.37 $\uparrow$)	51.11 51.52 (+0.41 $\uparrow$)	62.63 63.05 (+0.42 $\uparrow$)	39.59 40.32 (+0.73 $\uparrow$)
2/4	Adaround + CAT (Δ)	18.48 22.04 (+3.56 $\uparrow$)	–	0.36 0.74 (+0.38 $\uparrow$)	–	–	0.00 0.00 (0.00)
2/4	BRECQ + CAT (Δ)	62.69 62.67 (-0.02 $\downarrow$)	64.23 64.44 (+0.21 $\uparrow$)	0.12 0.20 (+0.08 $\uparrow$)	51.57 51.65 (+0.08 $\uparrow$)	64.39 64.57 (+0.18 $\uparrow$)	0.23 0.32 (+0.09 $\uparrow$)
2/4	Qdrop + CAT (Δ)	64.52 64.76 (+0.24 $\uparrow$)	69.95 70.50 (+0.55 $\uparrow$)	54.10 54.41 (+0.31 $\uparrow$)	63.17 63.50 (+0.33 $\uparrow$)	71.59 72.11 (+0.52 $\uparrow$)	63.39 64.09 (+0.70 $\uparrow$)
2/4	PD-Quant + CAT (Δ)	65.10 65.35 (+0.25 $\uparrow$)	70.87 71.29 (+0.42 $\uparrow$)	55.58 55.72 (+0.14 $\uparrow$)	63.90 64.31 (+0.41 $\uparrow$)	72.21 72.66 (+0.45 $\uparrow$)	63.06 63.82 (+0.76 $\uparrow$)
2/2	Adaround + CAT (Δ)	2.59 2.58 (-0.01 $\downarrow$)	–	–	–	–	–
2/2	BRECQ + CAT (Δ)	40.78 41.75 (+0.97 $\uparrow$)	13.78 16.69 (+2.91 $\uparrow$)	0.10 0.16 (+0.06 $\uparrow$)	0.68 2.03 (+1.35 $\uparrow$)	1.73 3.45 (+1.72 $\uparrow$)	0.12 0.18 (+0.06 $\uparrow$)
2/2	Qdrop + CAT (Δ)	51.47 51.76 (+0.29 $\uparrow$)	55.46 56.81 (+1.35 $\uparrow$)	11.81 12.25 (+0.44 $\uparrow$)	38.92 39.75 (+0.83 $\uparrow$)	54.12 55.26 (+1.14 $\uparrow$)	23.04 24.11 (+1.07 $\uparrow$)
2/2	PD-Quant + CAT (Δ)	52.86 53.18 (+0.32 $\uparrow$)	57.12 58.33 (+1.21 $\uparrow$)	13.92 14.33 (+0.41 $\uparrow$)	40.67 41.32 (+0.65 $\uparrow$)	55.23 56.24 (+1.01 $\uparrow$)	28.09 29.33 (+1.24 $\uparrow$)

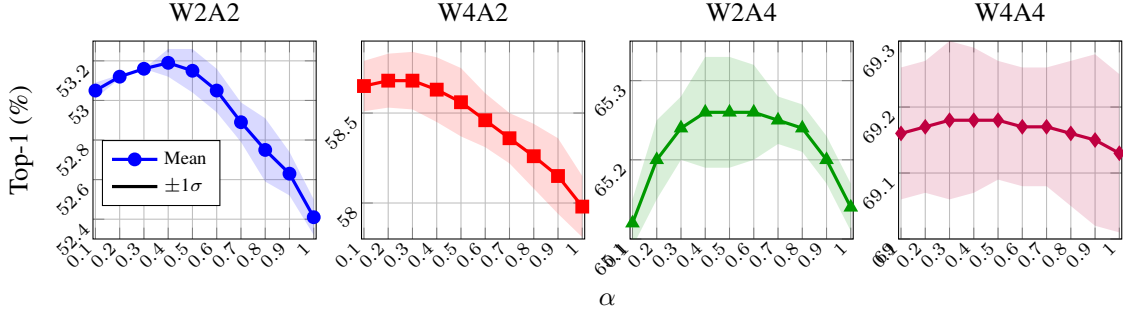


Figure 6: Ablation of α for ResNet-18 under different quantization precisions (W2A2, W2A4, W4A2, W4A4). The solid line denotes the mean Top-1 accuracy across three runs, and the shaded region indicates $\pm 1\sigma$. The comprehensive description of this ablation is provided in Section 5

D ABLATION OF α FOR DIFFERENT ARCHITECTURES

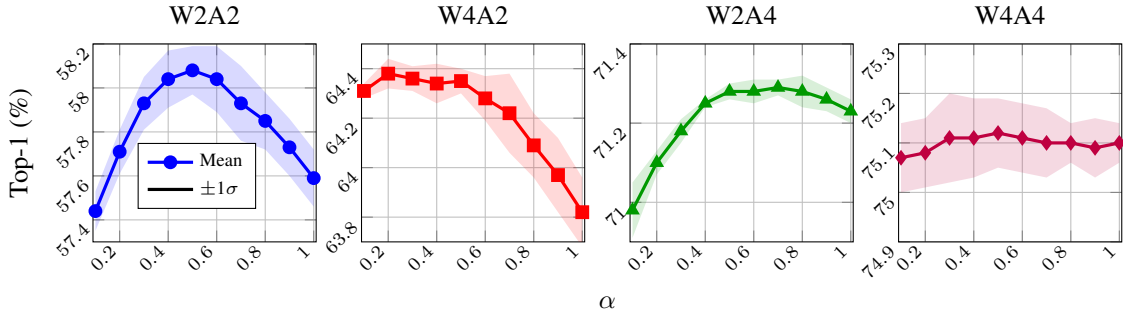


Figure 7: Ablation of α for ResNet-50 under different quantization precisions (W2A2, W2A4, W4A2, W4A4). The solid line denotes the mean Top-1 accuracy across three runs, and the shaded region indicates $\pm 1\sigma$. Across quantization settings, α controls the balance between the base quantized output and CAT correction. Performance peaks at moderate α ($\approx 0.3-0.5$) for W2A2 and W4A2, remains stable for W2A4, and is nearly flat for W4A4, indicating that CAT provides the most benefit in severely quantized regimes with low-bit activations.

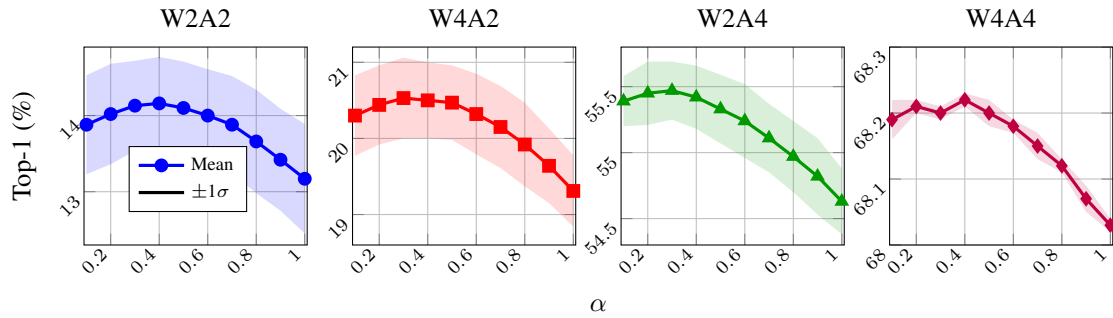


Figure 8: Ablation of α for MobileNetV2 under different quantization precisions (W2A2, W2A4, W4A2, W4A4). The solid line denotes the mean Top-1 accuracy across three runs, and the shaded region indicates $\pm 1\sigma$. For MobileNetV2, CAT gains are largest under severe activation quantization: W2A2 and W4A2 peak at moderate α (≈ 0.3 – 0.4) before dropping sharply, while W2A4 shows only mild variation and W4A4 remains essentially flat, indicating that α is most critical when activations are 2-bit.

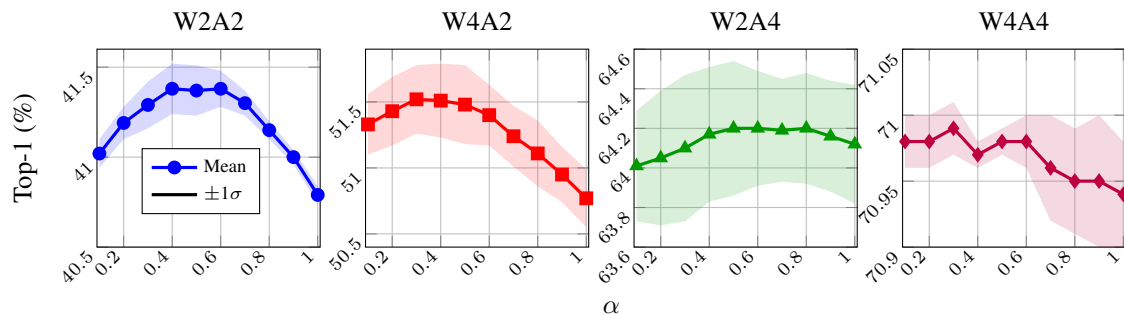


Figure 9: Ablation of α for RegNetX-600M under different quantization precisions (W2A2, W2A4, W4A2, W4A4). The solid line denotes the mean Top-1 accuracy across three runs, and the shaded region indicates $\pm 1\sigma$. For RegNetX-600M, α is most influential in 2-bit activation regimes: W2A2 and W4A2 peak at moderate α (≈ 0.3 – 0.5) and decline at larger values, whereas W2A4 shows only mild variation and W4A4 remains flat around 71%.

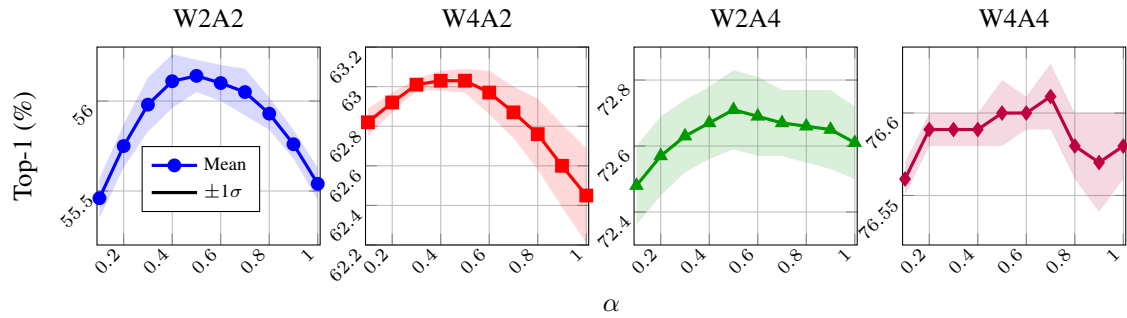


Figure 10: Ablation of α for RegNetX-3.2G under different quantization precisions (W2A2, W2A4, W4A2, W4A4). The solid line denotes the mean Top-1 accuracy across three runs, and the shaded region indicates $\pm 1\sigma$. For RegNetX-3.2G, α tuning is most beneficial in 2-bit activation regimes: W2A2 and W4A2 peak at moderate values ($\alpha \approx 0.3$ – 0.5) before declining, while W2A4 shows only mild variation and W4A4 remains flat near 76.6%.

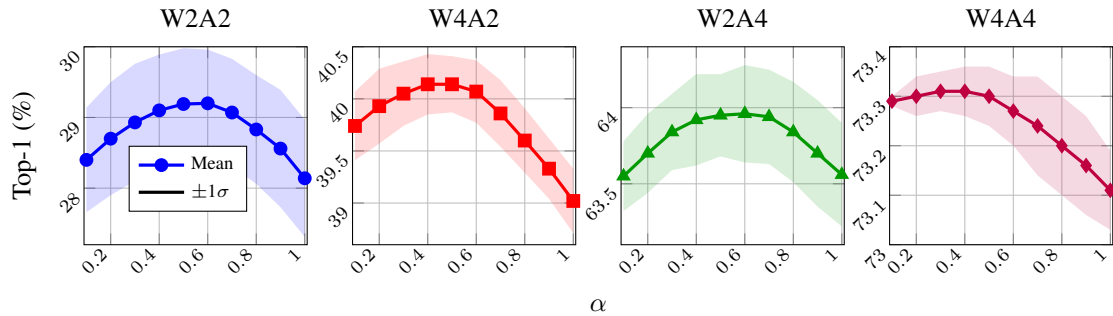


Figure 11: Ablation of α for MNasX2 under different quantization precisions (W2A2, W2A4, W4A2, W4A4). The solid line denotes the mean Top-1 accuracy across three runs, and the shaded region indicates $\pm 1\sigma$. For MNasX2, α tuning is most beneficial in 2-bit activation regimes: W2A2 and W4A2 peak at moderate values ($\alpha \approx 0.3$ – 0.5) before degrading, while W2A4 shows only mild variation and W4A4 remains flat around 73.3%.

E ABLATION OF THE NUMBER OF CLUSTERS FOR DIFFERENT ARCHITECTURES

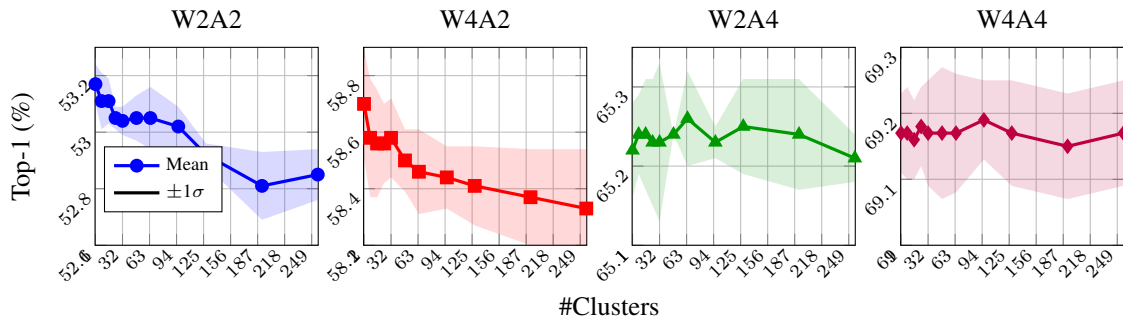


Figure 12: Ablation of the number of clusters for ResNet-18 (PCA=50, $\alpha = 0.6$) under different quantization precisions (W2A2, W4A2, W2A4; W4A4 pending). Solid line is the mean Top-1 over three runs; shaded region is $\pm 1\sigma$. The comprehensive description of this ablation is provided in Section 5

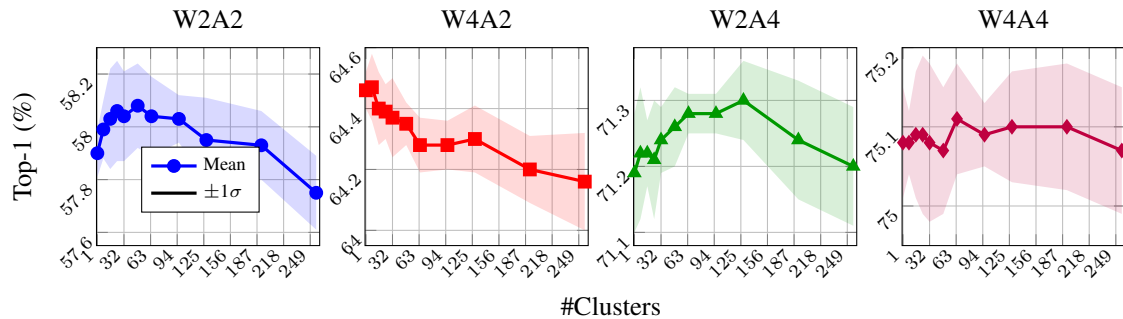


Figure 13: Ablation of number of clusters for ResNet-50 under different quantization precisions (W2A2, W2A4, W4A2, W4A4). The solid line denotes the mean Top-1 accuracy across three runs, and the shaded region indicates $\pm 1\sigma$. On ResNet-50, W2A2 benefits slightly from moderate K (peak $\sim 58.08\%$ near $K \approx 48$), whereas W4A2 declines with larger K , W2A4 shows a shallow peak near $K \approx 128$, and W4A4 is flat. Small-to-moderate K suffices; very large K offers no consistent gains.

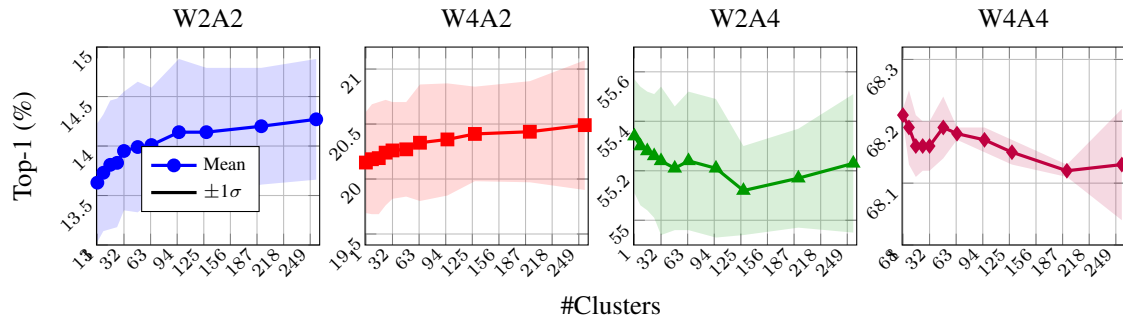


Figure 14: Ablation of number of clusters for MobileNetV2 under different quantization precisions (W2A2, W2A4, W4A2, W4A4). The solid line denotes the mean Top-1 accuracy across three runs, and the shaded region indicates $\pm 1\sigma$. On MobileNetV2, W2A2 and W4A2 improve steadily with larger K , peaking at $K=256$, while W2A4 declines slightly, and W4A4 is flat. Thus, more clusters help only in severe activation quantization.

F ABLATION OF PCA DIMENSIONS FOR DIFFERENT ARCHITECTURES

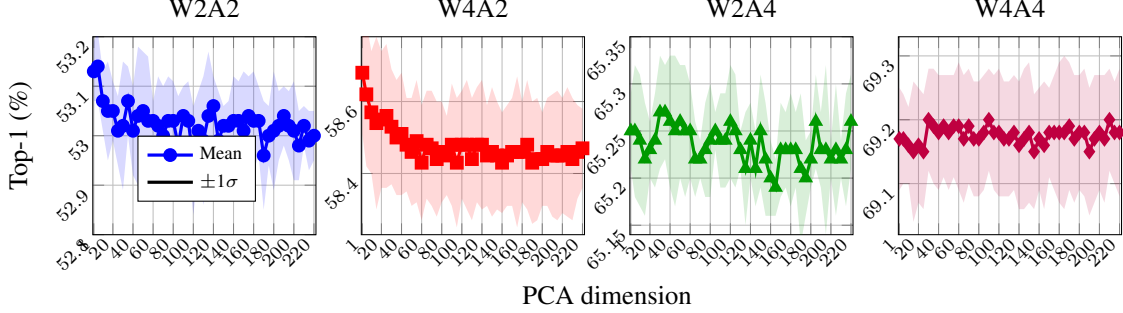


Figure 15: Ablation of PCA dimension k for ResNet-18 under different quantization precisions (W2A2, W4A2, W2A4, W4A4). Solid lines show mean Top-1 accuracy over three runs; the shaded region is $\pm 1\sigma$. The comprehensive description of this ablation is provided in Section 5

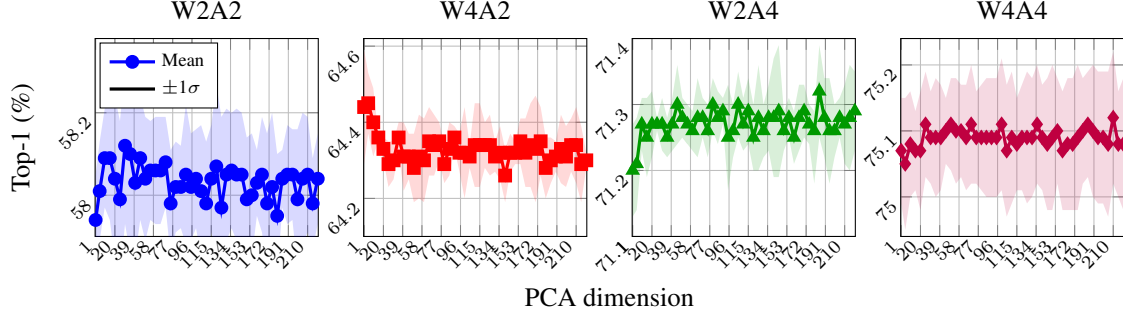


Figure 16: Ablation of PCA dimension d for ResNet-50 under different quantization precisions (W2A2, W2A4, W4A2, W4A4). The solid line is the mean Top-1 across three runs; the shaded band shows $\pm 1\sigma$. For ResNet-50, the number of clusters K controls how finely CAT partitions the logit space. Under 2-bit activations, performance improves rapidly as K increases from small values and peaks at a moderate range ($K \approx 8-32$), after which it plateaus or slightly declines due to over-partitioning and noisier parameter estimates. In W2A2, Top-1 rises steadily up to mid-range K and then saturates; W4A2 follows a similar pattern with a slightly earlier plateau. By contrast, W2A4 shows only mild gains across K , while W4A4 is essentially flat, reflecting the smaller FP-LQ gap in higher-precision regimes. Overall, these trends indicate that a moderate cluster count strikes the best balance between expressivity and robustness, especially when activations are quantized to 2-bit.

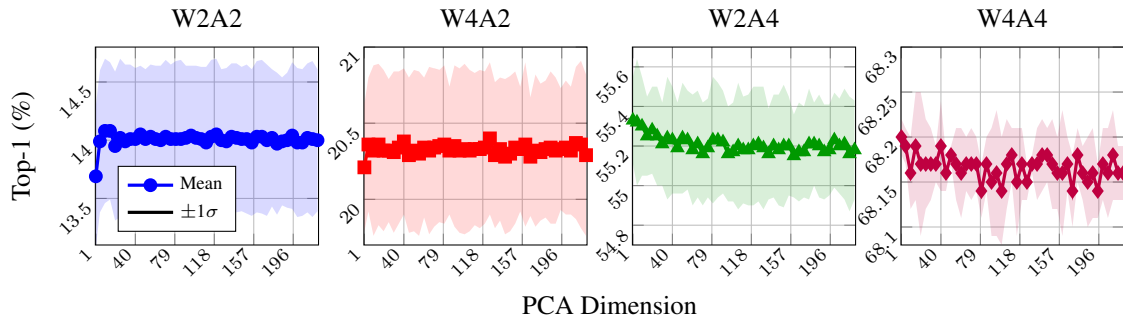


Figure 17: Effect of PCA dimension on **MobileNetV2** Top-1 accuracy under different quantization settings (W2A2, W4A2, W2A4, W4A4). The solid line is the mean over three runs; the shaded area is $\pm 1\sigma$. On MobileNetV2, PCA dimension matters mainly with 2-bit activations: W2A2 peaks at small p (~ 10 – 15) and then plateaus, while W4A2/W2A4 change little and W4A4 is flat. A compact subspace ($p \approx 10$ – 40) offers a good robustness–accuracy trade-off.

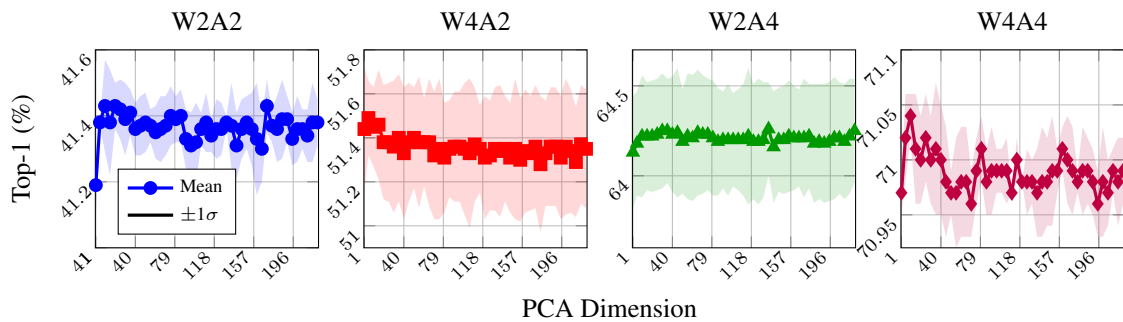


Figure 18: Effect of PCA dimension on **RegNetX-600MF** Top-1 accuracy under different quantization settings (W2A2, W4A2, W2A4, W4A4). The solid line is the mean over three runs; the shaded area is $\pm 1\sigma$. On RegNetX-600MF, PCA dimension matters only in W2A2, which peaks around $p=10$ – 20 and then plateaus; W4A2, W2A4, and W4A4 remain essentially flat. Thus, small subspaces suffice for this model.

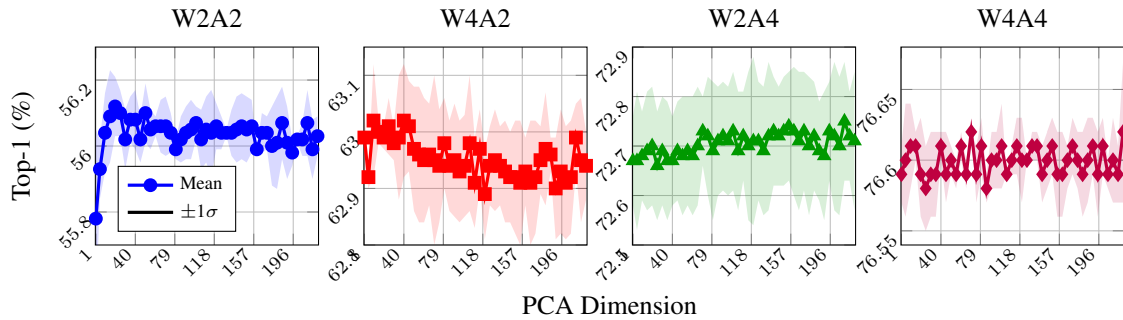


Figure 19: Effect of PCA dimension on **RegNetX-3.2GF** Top-1 accuracy under different quantization settings (W2A2, W4A2, W2A4, W4A4). The solid line is the mean over three runs; the shaded area is $\pm 1\sigma$. On RegNetX-3.2GF, W2A2 benefits slightly from small p (peaking near $p \approx 20$), while W4A2, W2A4, and W4A4 are essentially flat. A compact PCA subspace ($p \approx 10-40$) suffices for robust performance.

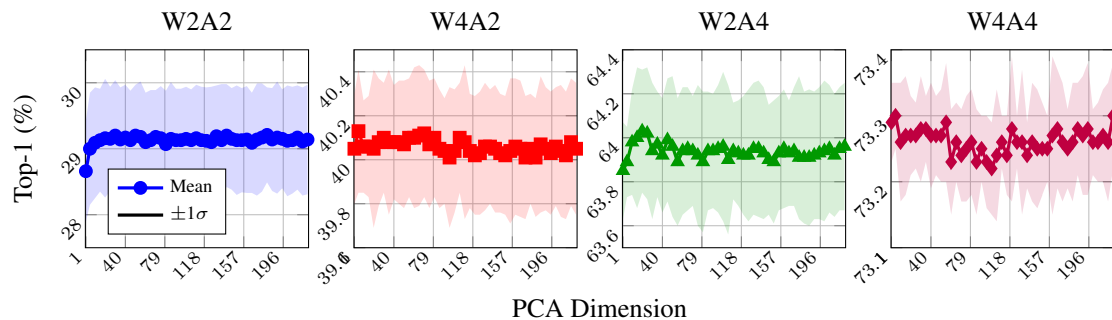


Figure 20: Effect of PCA dimension on **MNasX2** Top-1 accuracy under different quantization settings (W2A2, W4A2, W2A4, W4A4). The solid line is the mean over three runs; the shaded area is $\pm 1\sigma$. On MNasX2, W2A2 gains slightly up to $p \approx 30$ and then plateaus; W4A2, W2A4, and W4A4 are essentially flat. Thus, a compact PCA subspace ($p \approx 10-40$) suffices.

G ABLATION ON CAT FITTING SAMPLE SIZE

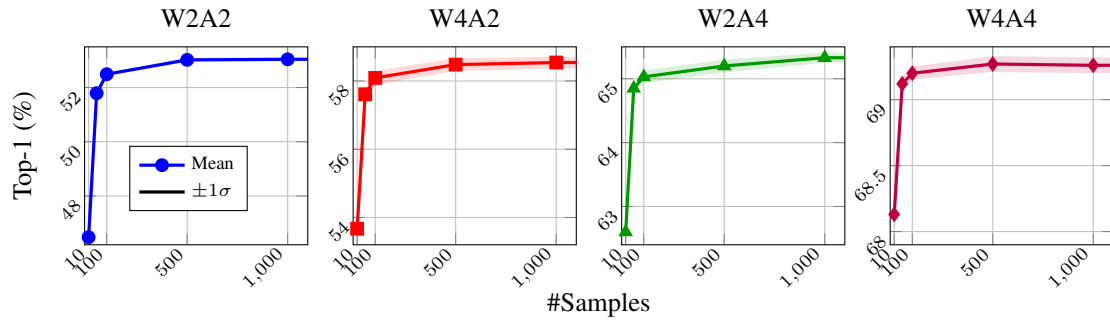


Figure 21: Ablation on the number of samples for ResNet-18 under different quantization settings (W2A2, W4A2, W2A4, W4A4). Solid lines show mean Top-1 accuracy; the shaded region is $\pm 1\sigma$. For ResNet-18, increasing the number of fitting samples yields clear gains that saturate once a few hundred to a thousand samples are used. In W2A2, Top-1 rises sharply from 46.5% at 10 samples to 53.0–53.1% by 500–1,000 samples, with only marginal changes thereafter (53.14–53.15% at 10k–100k). W4A2 follows the same trend, improving from 53.7% (10) to 58.5% (1k–10k) with diminishing returns beyond 500 samples. In W2A4, accuracy increases from 62.6% (10) to 65.3% (1k) and flattens around 65.2–65.3%. W4A4 similarly moves from 68.1% (10) to 69.3% (0.5k–10k) and then saturates. Overall, CAT achieves most of its benefit with ~ 500 –1,000 samples, indicating that modest data budgets suffice for effective post-training error restoration.

H CAT IMPROVEMENT ON ViT QUANTIZATION VIA DIFFERENT PTQ METHODS

Table 6: Top-1 accuracy (%) for BRECC and QDrop with/without CAT under W2A2 and W4A4.

Method	Variant	W2A2						W4A4					
		ViT-S	ViT-B	Swin-S	Swin-B	DeiT-S	DeiT-B	ViT-S	ViT-B	Swin-S	Swin-B	DeiT-S	DeiT-B
BRECC	Base	0.63	2.042	3.43	4.1	4.84	12.46	73.96	79.23	81.07	82.6	69.82	77.64
	+ CAT	0.75	2.17	3.52	4.27	5.04	12.79	74.09	79.42	81.11	82.65	70.09	77.7
QDrop	Base	3.36	4.77	12.85	15.69	4.1	30.53	69.51	81.96	80.79	82.25	69.93	78.01
	+ CAT	3.34	4.83	12.83	15.77	4.27	30.72	69.54	82	80.86	82.30	70.13	78.05

Table 6 reports Top-1 accuracy for ViT, Swin, and DeiT models under W2A2 and W4A4 quantization with BRECC and QDrop. CAT consistently improves performance across all settings. The gains are most notable in ultra-low precision (W2A2), with improvements up to +0.3%–0.4%, while W4A4 shows smaller but steady benefits, confirming CAT’s generalizability across architectures and quantization methods. Table 7

Table 7: Top-1 accuracy (%) for BRECC and QDrop with/without CAT on DeiT-S and DeiT-B under different precisions.

Method	Variant	W2A2		W4A4		W2A6		W4A6		W6A6	
		DeiT-S	DeiT-B	DeiT-S	DeiT-B	DeiT-S	DeiT-B	DeiT-S	DeiT-B	DeiT-S	DeiT-B
BRECC	Base	4.84	12.46	69.82	77.64	65.25	75.22	77.64	80.83	76.7	80.49
	+ CAT	5.04	12.79	70.09	77.70	65.30	75.28	77.65	80.9	76.87	80.5
QDrop	Base	4.10	30.53	69.93	78.01	65.71	76.0	77.9	80.9	77.74	80.82
	+ CAT	4.27	30.72	70.13	78.05	66.04	76.11	77.95	80.94	77.82	80.86

shows that CAT consistently enhances accuracy for both DeiT-S and DeiT-B across all tested precisions. The improvements are most pronounced under very low-bit precision, reaching up to +0.4% at W2A2 (e.g., DeiT-B with BRECC: 12.79% vs. 12.46%) and +0.3% at W2A6 (e.g., DeiT-S with QDrop: 66.04% vs. 65.71%). At more moderate bit-widths (W4A4, W4A6, W6A6), CAT provides smaller but consistent gains in the range of +0.1–0.2%, indicating that the benefit decreases as the quantized model approaches full-precision accuracy. Importantly, improvements are observed for both BRECC and QDrop baselines and across both DeiT-S and DeiT-B models, highlighting CAT’s robustness and generality as a lightweight correction mechanism for transformer architectures under different quantization regimes.

I DISCUSSION

Our experiments demonstrate that CAT consistently improves the performance of a wide range of PTQ baselines across architectures and precision settings. On average, we observe absolute accuracy gains of +1.56% for W2A2, +1.19% for W4A2, +0.39% for W2A4, and +0.10% for W4A4 within the largest parameter regime (20M – 40M). Smaller-capacity models (< 10M) also benefit, though to a lesser extent, with gains of +0.68% (W2A2) and +0.96% (W4A2). These trends confirm that CAT is particularly effective in challenging low-precision regimes, while improvements diminish as precision increases. A key observation is that lower activation precision yields the largest relative improvements. In both W2A2 and W4A2 regimes, CAT provides the strongest corrections, while W2A4 and W4A4 exhibit smaller yet consistent gains. This pattern indicates that activation quantization is the dominant factor driving representational degradation, and CAT’s cluster-specific corrections directly mitigate this bottleneck. In contrast, weight precision plays a comparatively minor role: lowering weights to 2-bit (while keeping activations higher) causes less severe degradation and thus smaller relative improvements with CAT.

We also find that larger models ($20M - 40M$ parameters, e.g., ResNet-50, RegNetX-3.2GF) exhibit the highest absolute improvements in LQ settings, surpassing $+1\%$ in both W2A2 and W4A2. Our interpretation is that these models possess greater capacity to process complex information, but when activations are severely quantized, their representational potential is limited by activation precision. CAT effectively restores this information flow by aligning low-precision activations with their full-precision distributions. In contrast, compact models ($3-7M$, e.g., MobileNetV2, MNasX2) still benefit from CAT but show smaller relative gains, as their limited redundancy constrains the extent to which CAT can compensate. In summary, the results highlight three consistent trends: (i) CAT is most effective when activation precision is low, (ii) larger models benefit more in absolute terms, and (iii) weight quantization alone has a weaker influence on CAT’s efficacy compared to activation quantization. Together, these findings position CAT as a practical and robust enhancement for PTQ pipelines, particularly in ultra-low-precision and high-capacity scenarios where quantization errors are most detrimental.

Table 8: Average Top-1 improvement Δ (%) of CAT over Base by parameter regimes and precision. Regimes: $< 10M$ (MobileNetV2, MNasX2, RegNetX-600MF), $10M - 20M$ (ResNet-18), $20M - 40M$ (ResNet-50, RegNetX-3.2GF). Darker green indicates larger improvement.

Params Regime	W2A2	W4A2	W2A4	W4A4
$< 10M$	0.68	0.96	0.32	0.37
$10M - 20M$	0.53	0.47	0.16	0.07
$20M - 40M$	1.56	1.19	0.39	0.10