

QUANT-DLLM: POST-TRAINING EXTREME LOW-BIT QUANTIZATION FOR DIFFUSION LARGE LANGUAGE MODELS

Anonymous authors

Paper under double-blind review

CONTENTS

A First-order DAQ	2
B K-order DAQ	3
B.1 Algorithmic Process of K-order DAQ	3
C Justification for the Data-aware Frobenius Proxy	6
C.1 The Challenge: Parameter Coupling in RSR	6
C.2 The Solution: Data-aware Objective Reformulation (DOR)	6
D Masked-diffusion decoding	6
E Detailed Experiments Settings	7
E.1 Settings for LLaDA Models	7
E.2 Settings for Dream Models	7
F More Experiments	8
F.1 Effectiveness of MCS on Baseline Methods	8

A FIRST-ORDER DAQ

The First-order Data-aware Any-order Quantizer (DAQ) aims to find a 1-bit approximation of a weight matrix by minimizing a quantization error that is weighted by the importance of each parameter. This method acknowledges that not all weights contribute equally to the model’s performance and thus prioritizes the fidelity of more salient weights during quantization.

The objective function, consistent with the main text, is to minimize the weighted Frobenius norm of the quantization error:

$$\mathcal{L} = \|\Lambda \odot (\mathbf{W} - \widehat{\mathbf{W}})\|_F^2 \quad (1)$$

where $\mathbf{W} \in \mathbb{R}^{n \times m}$ is the full-precision weight matrix, $\Lambda \in \mathbb{R}^{n \times m}$ is the importance mask derived from calibration data, and $\odot$ denotes the Hadamard (element-wise) product. The quantized weight matrix, $\widehat{\mathbf{W}}$, is parameterized as follows:

$$\widehat{\mathbf{W}} = (\alpha^r (\alpha^c)^\top) \odot \mathbf{B} \quad (2)$$

Here, $\alpha^r \in \mathbb{R}^n$ and $\alpha^c \in \mathbb{R}^m$ are the row and column scaling factor vectors, respectively, and $\mathbf{B} \in \{-1, +1\}^{n \times m}$ is the binary matrix.

The complete algorithm for the first-order DAQ is detailed in Algorithm 1. The process consists of three main stages: initialization, iterative refinement of scaling factors, and an optional re-rounding of the binary matrix.

1. Initialization The algorithm begins by establishing an initial approximation of the full-precision weights $\mathbf{W}$. This involves initializing the binary matrix $\mathbf{B}$ and the scaling factors α^r and α^c .

First, the binary matrix $\mathbf{B}$ is determined by the sign of the original weights, capturing their fundamental polarity:

$$\mathbf{B} = \text{sign}(\mathbf{W}) \quad (3)$$

Next, inspired by ARB-LLM, the row and column scaling factors are initialized to approximate the magnitude of the weights. The row scaling factor α^r is initialized as the mean absolute value of each row:

$$\alpha^r = \frac{1}{m} |\mathbf{W}| \mathbf{1}_m \quad (4)$$

where $|\mathbf{W}|$ is the element-wise absolute value of $\mathbf{W}$ and $\mathbf{1}_m$ is an all-ones vector of size m . Subsequently, the column scaling factor α^c is initialized based on the row-scaled weights:

$$\alpha^c = \frac{1}{n} |\mathbf{W}|^\top (\alpha^r)^{\circ-1} \quad (5)$$

where $(\alpha^r)^{\circ-1}$ denotes the element-wise inverse of α^r . This initialization provides a reasonable starting point for the subsequent iterative optimization.

2. RSR Following initialization, the algorithm enters an iterative loop to refine the scaling factors α^r and α^c for a total of T rounds using the Row-column Successive Re-scaling (RSR) procedure. Within each iteration, we alternately update the scaling factors. Let $\Lambda^2 = \Lambda \odot \Lambda$ be the element-wise square of the importance mask.

The *update_* α^r function solves $\nabla_{\alpha^r} \mathcal{L} = \mathbf{0}$ and updates the entire α^r vector in closed form:

$$\alpha^r \leftarrow [(\Lambda^2 \odot \mathbf{W} \odot \mathbf{B}) \alpha^c] \oslash [\Lambda^2 (\alpha^c \odot \alpha^c) + \varepsilon \mathbf{1}_n] \quad (6)$$

where $\oslash$ denotes element-wise division and ε is a small constant for numerical stability.

Similarly, the *update_* α^c function solves $\nabla_{\alpha^c} \mathcal{L} = \mathbf{0}$ using the newly updated α^r :

$$\alpha^c \leftarrow [(\Lambda^2 \odot \mathbf{W} \odot \mathbf{B})^\top \alpha^r] \oslash [(\Lambda^2)^\top (\alpha^r \odot \alpha^r) + \varepsilon \mathbf{1}_m] \quad (7)$$

This alternating update process monotonically decreases the data-aware objective function $\mathcal{L}$, progressively improving the approximation of the quantized matrix $\widehat{\mathbf{W}}$.

3. Optional Re-rounding of $\mathbf{B}$ After updating the scaling factors in an iteration, the binary matrix $\mathbf{B}$ can be optionally re-evaluated. While holding $\mathbf{B}$ fixed during iterations simplifies optimization,

Algorithm 1 DAQ¹: First-order Data-aware Quantizer

```

108 func DAQ1(W, Z,  $T$ ,  $\lambda$ ,  $\varepsilon$ )
109 Input  $\mathbf{W} \in \mathbb{R}^{n \times m}$ ,  $\mathbf{Z} \in \mathbb{R}^{n \times m}$ ,  $T \in \mathbb{N}$ ,  $\lambda > 1$ ,
110  $\varepsilon > 0$ 
111 Output  $\hat{\mathbf{W}} \in \mathbb{R}^{n \times m}$ ,  $\alpha_r \in \mathbb{R}^n$ ,  $\alpha_c \in \mathbb{R}^m$ ,
112  $\mathbf{B} \in \{\pm 1\}^{n \times m}$ 
113   1:  $\mathbf{\Lambda} := \text{build\_importance\_mask}(\mathbf{Z}, \lambda)$ 
114   2:  $\alpha_r, \alpha_c, \mathbf{B} := \text{binary\_rc\_init}(\mathbf{W})$ 
115   3: for  $t = 1, 2, \dots, T$  do
116     4:  $\alpha_r \leftarrow \text{update\_}\alpha^r(\mathbf{W}, \mathbf{B}, \alpha_c, \mathbf{\Lambda}, \varepsilon)$ 
117     5:  $\alpha_c \leftarrow \text{update\_}\alpha^c(\mathbf{W}, \mathbf{B}, \alpha_r, \mathbf{\Lambda}, \varepsilon)$ 
118     6: optional  $\mathbf{B} \leftarrow \text{sign}(\mathbf{W})$ 
119   7: end for
120   8:  $\hat{\mathbf{W}} := (\alpha_r \alpha_c^\top) \odot \mathbf{B}$ 
121   9: return  $\hat{\mathbf{W}}, \alpha_r, \alpha_c, \mathbf{B}$ 
122
123 func binary_rc_init(W)
124   1:  $\mathbf{B} := \text{sign}(\mathbf{W})$ 
125   2:  $\alpha_r := \frac{1}{m} |\mathbf{W}| \mathbf{1}_m$ 
126   3:  $\alpha_c := \frac{1}{n} |\mathbf{W}|^\top (\alpha_r)^{\circ -1}$ 
127   4: return  $\alpha_r, \alpha_c, \mathbf{B}$ 
128
129 func update_ $\alpha^r(\mathbf{X}, \mathbf{B}, \alpha_c, \mathbf{\Lambda}, \varepsilon)$ 
130   1:  $\mathbf{\Lambda}^{(2)} := \mathbf{\Lambda} \odot \mathbf{\Lambda}$ 
131   2:  $\mathbf{u} := (\mathbf{\Lambda}^{(2)} \odot \mathbf{X} \odot \mathbf{B}) \alpha_c$ 
132   3:  $\mathbf{v} := \mathbf{\Lambda}^{(2)} (\alpha_c^{\circ 2}) + \varepsilon \mathbf{1}_n$ 
133   4: return  $\mathbf{u} \oslash \mathbf{v}$ 
134
135 func update_ $\alpha^c(\mathbf{X}, \mathbf{B}, \alpha_r, \mathbf{\Lambda}, \varepsilon)$ 
136   1:  $\mathbf{\Lambda}^{(2)} := \mathbf{\Lambda} \odot \mathbf{\Lambda}$ 
137   2:  $\mathbf{u} := (\mathbf{\Lambda}^{(2)} \odot \mathbf{X} \odot \mathbf{B})^\top \alpha_r$ 
138   3:  $\mathbf{v} := (\mathbf{\Lambda}^{(2)})^\top (\alpha_r^{\circ 2}) + \varepsilon \mathbf{1}_m$ 
139   4: return  $\mathbf{u} \oslash \mathbf{v}$ 
140
141 func build_importance_mask(Z,  $\lambda$ )
142   1:  $\mu := \text{mean}(\mathbf{Z})$ ,  $\sigma := \text{std}(\mathbf{Z})$ 
143   2:  $\mathbf{M} := \mathbb{I}(|\mathbf{Z} - \mu| > 3\sigma)$ 
144   3: return  $\mathbf{1}_{n \times m} + (\lambda - 1) \mathbf{M}$ 

```

re-rounding it can sometimes further reduce quantization error. The optimal binary matrix that minimizes the objective is determined by the sign of the weights:

$$\mathbf{B} \leftarrow \text{sign}(\mathbf{W}) \quad (8)$$

In practice, for the first-order case, this optional step simply re-aligns the binary matrix with the original weights, reinforcing the initial structure. This is particularly useful if the iterative scaling process causes a significant shift in the reconstructed magnitudes.

B K-ORDER DAQ

To further enhance the representational capacity of the quantizer under a fixed bit-budget, the DAQ framework is extended to a K-order representation. This approach approximates the full-precision weight matrix $\mathbf{W}$ as a sum of K distinct first-order components. Each component consists of its own binary matrix $\mathbf{B}^{(k)}$ and a corresponding pair of row-column scaling factors, $\alpha_r^{(k)}$ and $\alpha_c^{(k)}$. The K-order quantized weight matrix $\hat{\mathbf{W}}$ is formulated as:

$$\hat{\mathbf{W}} = \sum_{k=1}^K \left(\alpha_r^{(k)} (\alpha_c^{(k)})^\top \right) \odot \mathbf{B}^{(k)} \quad (9)$$

The optimization for this K-order representation is performed through a greedy, iterative process, as detailed in Algorithm 2. The procedure can be broken down into three primary stages: greedy initialization, iterative refinement of all components, and a final update of the binary matrices.

B.1 ALGORITHMIC PROCESS OF K-ORDER DAQ

1. Greedy Initialization The algorithm begins with a greedy, stage-wise construction of the K components. This process iteratively fits a new first-order component to the residual error from the previous stages. For $k = 1, 2, \dots, K$:

1. **Compute Current Residual Matrix:** The residual matrix $\mathbf{R}^{(k)}$ is the difference between the original weights and the sum of all previously initialized components.

$$\mathbf{R}^{(k)} = \mathbf{W} - \sum_{q=1}^{k-1} \left(\left(\alpha_r^{(q)} (\alpha_c^{(q)})^\top \right) \odot \mathbf{B}^{(q)} \right) \quad (10)$$

For the first iteration ($k = 1$), the residual is the full-precision weight matrix itself, $\mathbf{R}^{(1)} = \mathbf{W}$.

2. **Initialize the k -th Component:** The *binary_rc_init* function is called on the current residual to obtain the initial parameters for the k -th component.

$$\alpha_r^{(k)}, \alpha_c^{(k)}, \mathbf{B}^{(k)} := \text{binary_rc_init}(\mathbf{R}^{(k)}) \quad (11)$$

The internal operations of this function are:

$$\mathbf{B}^{(k)} = \text{sign}(\mathbf{R}^{(k)}) \quad (12)$$

$$\alpha_r^{(k)} = \frac{1}{m} |\mathbf{R}^{(k)}| \mathbf{1}_m \quad (13)$$

$$\alpha_c^{(k)} = \frac{1}{n} |\mathbf{R}^{(k)}|^\top (\alpha_r^{(k)})^{\circ -1} \quad (14)$$

This greedy procedure provides a strong starting point for the joint optimization phase.

2. Joint RSR After initializing all K components, the algorithm enters a main iterative loop to jointly refine all parameters for T rounds. In each round, the algorithm cycles through every component $k = 1, \dots, K$ and refines its scaling factors while keeping all other components fixed. For each component k :

1. **Compute Partial Residual Matrix:** The partial residual $\mathbf{R}^{(k)}$ is calculated by subtracting all components *except* the k -th one from the original weights. This isolates the target for the current component's optimization.

$$\mathbf{R}^{(k)} = \mathbf{W} - \sum_{q \neq k} \left((\alpha_r^{(q)} (\alpha_c^{(q)})^\top) \odot \mathbf{B}^{(q)} \right) \quad (15)$$

2. **Update Row and Column Scaling Factors:** The RSR procedure is applied to find the optimal scaling factors for the k -th component that best approximate the partial residual $\mathbf{R}^{(k)}$. Let $\Lambda^2 = \Lambda \odot \Lambda$.

$$\alpha_r^{(k)} \leftarrow \left[(\Lambda^2 \odot \mathbf{R}^{(k)} \odot \mathbf{B}^{(k)}) \alpha_c^{(k)} \right] \oslash \left[\Lambda^2 (\alpha_c^{(k)} \odot \alpha_c^{(k)}) + \varepsilon \mathbf{1}_n \right] \quad (16)$$

$$\alpha_c^{(k)} \leftarrow \left[(\Lambda^2 \odot \mathbf{R}^{(k)} \odot \mathbf{B}^{(k)})^\top \alpha_r^{(k)} \right] \oslash \left[(\Lambda^2)^\top (\alpha_r^{(k)} \odot \alpha_r^{(k)}) + \varepsilon \mathbf{1}_m \right] \quad (17)$$

This process allows the parameters of all components to be jointly optimized to minimize the overall quantization error $\mathcal{L}$.

3. Update of Binary Matrices After the iterative refinement of the scaling factors within a round, the binary matrices $\{\mathbf{B}^{(k)}\}_{k=1}^K$ are jointly updated. Since the binary values are discrete, this is solved via an exhaustive search for each weight position (i, j) .

1. **Construct Scale Vector:** For each position (i, j) , a vector $\mathbf{s}_{ij} \in \mathbb{R}^K$ is constructed, containing the scaling products from all K components.

$$\mathbf{s}_{ij} = \begin{bmatrix} (\alpha_{r,i}^{(1)})(\alpha_{c,j}^{(1)}) \\ \vdots \\ (\alpha_{r,i}^{(K)})(\alpha_{c,j}^{(K)}) \end{bmatrix} \quad (18)$$

2. **Minimal-Error Search:** The *update_B* function finds the optimal binary vector $\mathbf{b}_{ij} \in \{-1, +1\}^K$ that minimizes the reconstruction error for W_{ij} .

$$\mathbf{b}_{ij} = \arg \min_{\mathbf{b} \in \{-1, +1\}^K} |W_{ij} - \mathbf{s}_{ij}^\top \mathbf{b}| \quad (19)$$

The elements of the resulting vector $\mathbf{b}_{ij} = [b_1, \dots, b_K]^\top$ are then used to update the corresponding entries in the binary matrices:

$$(\mathbf{B}^{(k)})_{ij} \leftarrow b_k, \quad \text{for } k = 1, \dots, K \quad (20)$$

This step ensures that the binary carriers are optimally aligned with the refined scaling factors after each round of optimization.

Algorithm 2 DAQ: Data-aware any-order Quantizer

```

func DAQ( $\mathbf{W}, \mathbf{Z}, K, T, \lambda$ )
Input  $\mathbf{W}, \mathbf{Z} \in \mathbb{R}^{n \times m}$ ,  $K, T \in \mathbb{N}$ ,  $\lambda > 1$ 
Output
 $\hat{\mathbf{W}} = \sum_{k=1}^K (\alpha_r^{(k)} \alpha_c^{(k)\top}) \odot \mathbf{B}^{(k)}$ 
1:  $\Lambda := \text{build\_importance\_mask}(\mathbf{Z}, \lambda)$ 
2:  $\hat{\mathbf{W}} := \mathbf{0}_{n \times m}$ 
3: for  $k = 1, 2, \dots, K$  do
4:    $\mathbf{R} \leftarrow \mathbf{W} - \hat{\mathbf{W}}$ 
5:    $\alpha_r^{(k)}, \alpha_c^{(k)}, \mathbf{B}^{(k)} \leftarrow \text{binary\_rc\_init}(\mathbf{R})$ 
6:    $\mathbf{S}^{(k)} \leftarrow \alpha_r^{(k)} (\alpha_c^{(k)})^\top$ 
7:    $\hat{\mathbf{W}} \leftarrow \hat{\mathbf{W}} + \mathbf{S}^{(k)} \odot \mathbf{B}^{(k)}$ 
8: end for
9: for  $t = 1, 2, \dots, T$  do
10:  for  $k = 1, 2, \dots, K$  do
11:     $\hat{\mathbf{W}}_{\setminus k} \leftarrow \sum_{q \neq k} (\mathbf{S}^{(q)} \odot \mathbf{B}^{(q)})$ 
12:     $\mathbf{R}^{(k)} \leftarrow \mathbf{W} - \hat{\mathbf{W}}_{\setminus k}$ 
13:     $\alpha_r^{(k)} \leftarrow \text{update\_}\alpha_r(\mathbf{R}^{(k)}, \mathbf{B}^{(k)}, \alpha_c^{(k)}, \Lambda)$ 
14:     $\alpha_c^{(k)} \leftarrow \text{update\_}\alpha_c(\mathbf{R}^{(k)}, \mathbf{B}^{(k)}, \alpha_r^{(k)}, \Lambda)$ 
15:     $\mathbf{S}^{(k)} \leftarrow \alpha_r^{(k)} (\alpha_c^{(k)})^\top$ 
16:  end for
17:   $\{\mathbf{B}^{(k)}\}_{k=1}^K \leftarrow \text{update\_B}(\mathbf{W}, \{\alpha_r^{(k)}\}_{k=1}^K, \{\alpha_c^{(k)}\}_{k=1}^K)$ 
18:   $\hat{\mathbf{W}} \leftarrow \sum_{k=1}^K \mathbf{S}^{(k)} \odot \mathbf{B}^{(k)}$ 
19: end for
20: return  $\hat{\mathbf{W}}$ 

func binary_rc_init( $\mathbf{X}$ )
1:  $\alpha_r \leftarrow \frac{1}{m} |\mathbf{X}| \mathbf{1}_m$ 
2:  $\alpha_c \leftarrow \frac{1}{n} |\mathbf{X}|^\top \text{diag}(\alpha_r)^{-1} \mathbf{1}_n$ 
3:  $\mathbf{B} \leftarrow \text{sign}(\mathbf{X})$ 
4: return  $\alpha_r, \alpha_c, \mathbf{B}$ 

func update_αr( $\mathbf{X}, \mathbf{B}, \alpha_c, \Lambda$ )
1:  $\mathbf{u} \leftarrow (\Lambda^2 \odot \mathbf{X} \odot \mathbf{B}) \alpha_c$ 
2:  $\mathbf{v} \leftarrow (\Lambda^2 (\text{diag} \alpha_c) \alpha_c) + \varepsilon \mathbf{1}_n$ 
3: return  $\mathbf{u} \oslash \mathbf{v}$ 

func update_αc( $\mathbf{X}, \mathbf{B}, \alpha_r, \Lambda$ )
1:  $\mathbf{u} \leftarrow (\Lambda^2 \odot \mathbf{X} \odot \mathbf{B})^\top \alpha_r$ 
2:  $\mathbf{v} \leftarrow (\Lambda^2)^\top (\text{diag}(\alpha_r) \alpha_r) + \varepsilon \mathbf{1}_m$ 
3: return  $\mathbf{u} \oslash \mathbf{v}$ 

func update_B( $\mathbf{W}, \{\alpha_r^{(k)}\}_{k=1}^K, \{\alpha_c^{(k)}\}_{k=1}^K$ )
1: for  $i = 1, \dots, n$  do
2:   for  $j = 1, \dots, m$  do
3:      $\mathbf{s} \leftarrow [(\alpha_r^{(k)})_i (\alpha_c^{(k)})_j]_{k=1}^K$ 
4:      $\mathbf{b} \leftarrow \text{search}(W_{ij}, \mathbf{s})$ 
5:     for  $k = 1, \dots, K$  do
6:        $(\mathbf{B}^{(k)})_{ij} \leftarrow \mathbf{b}_k$ 
7:     end for
8:   end for
9: end for
10: return  $\{\mathbf{B}^{(k)}\}_{k=1}^K$ 

func build_importance_mask( $\mathbf{Z}, \lambda$ )
1:  $\mu \leftarrow \text{mean}(\mathbf{Z}), \sigma \leftarrow \text{std}(\mathbf{Z})$ 
2:  $\mathbf{M} \leftarrow \mathbb{I}(|\mathbf{Z} - \mu| > 3\sigma)$ 
3: return  $\mathbf{1}_{n \times m} + (\lambda - 1) \mathbf{M}$ 

```

C JUSTIFICATION FOR THE DATA-AWARE FROBENIUS PROXY

In our Quant-dLLM framework, the ideal objective is to directly minimize the quantization error on the layer’s output. This is captured by the data-aware objective $\mathcal{L}_2$, which utilizes the second-moment matrix $\mathbf{S}_{\text{MCS}}$ computed from the calibration data generated by our Masked Calibration Simulation (MCS) method:

$$\mathcal{L}_2 = \text{Tr}((\mathbf{W} - \widehat{\mathbf{W}}) \mathbf{S}_{\text{MCS}} (\mathbf{W} - \widehat{\mathbf{W}})^\top) \quad (21)$$

While this objective is theoretically superior, its direct optimization with our separable row-column (RC) scaling, $\widehat{\mathbf{W}} = (\alpha^r (\alpha^c)^\top) \odot \mathbf{B}$, is computationally intractable. This section justifies our use of a tractable proxy.

C.1 THE CHALLENGE: PARAMETER COUPLING IN RSR

The primary difficulty arises from **parameter coupling** when attempting to derive an update rule for the column scaling factors α^c within our Row-column Successive Re-scaling (RSR) algorithm. When we take the partial derivative of $\mathcal{L}_2$ with respect to a single column scaling factor, $\alpha_{c,t}$, and set it to zero, the resulting equation for $\alpha_{c,t}$ remains dependent on all other column scaling factors $\alpha_{c,k}$ (where $k \neq t$).

This coupling is illustrated by the conceptual structure of the resulting equation:

$$\frac{\partial \mathcal{L}_2}{\partial \alpha_{c,t}} \propto \sum_{i=1}^n \alpha_{r,i}^2 B_{it}^2 \sum_{k=1}^m S_{tk} \alpha_{c,k} - \sum_{i=1}^n \alpha_{r,i} B_{it} \sum_{k=1}^m W_{ik} S_{kt} = 0 \quad (22)$$

The term $\sum_{k=1}^m S_{tk} \alpha_{c,k}$ clearly demonstrates the problem: due to the non-diagonal nature of the $\mathbf{S}_{\text{MCS}}$ matrix (where S_{tk} can be non-zero for $t \neq k$), the optimal value for $\alpha_{c,t}$ is part of a system of linear equations involving the entire vector α^c . Solving this dense system at each step is computationally prohibitive for a PTQ method.

C.2 THE SOLUTION: DATA-AWARE OBJECTIVE REFORMULATION (DOR)

To overcome this challenge, we introduced the **Data-aware Objective Reformulation (DOR)** in the main paper. DOR serves as a bridge, creating a computationally efficient yet effective proxy for the intractable $\mathcal{L}_2$ objective. By analyzing the importance matrix $\mathbf{Z}$ derived from $\mathbf{S}_{\text{MCS}}$, DOR constructs a sparse importance mask Λ .

This allows us to optimize the **data-aware Frobenius proxy**:

$$\widehat{\mathcal{L}}_\Lambda(\alpha_r, \alpha_c) = \|\Lambda \odot (\mathbf{W} - \widehat{\mathbf{W}})\|_F^2 \quad (23)$$

This objective successfully retains the critical data-dependent information (by up-weighting salient elements) while being structurally suitable for our efficient, closed-form RSR updates.

D MASKED-DIFFUSION DECODING

Gaussian masked diffusion (MDM). For a continuous input x and a binary mask $r \in \{0, 1\}^L$, we perturb only the masked coordinates to obtain x_t and train a noise predictor $\varepsilon_\theta(x_t, r, t)$ with the standard denoising objective

$$\mathcal{L}_{\text{MDM}} = \mathbb{E}_{x, r, \varepsilon, t} \|\varepsilon - \varepsilon_\theta(x_t, r, t)\|_2^2, \quad \varepsilon \sim \mathcal{N}(0, I). \quad (24)$$

Absorbing-state discrete variant. For token sequences $y \in \{e_1, \dots, e_K\}^L$ with absorbing token M (one-hot m) and a monotone visibility schedule $\alpha_t \in [0, 1]$ (we use a linear schedule by default), the forward corruption is

$$q(y_t | y) = \text{Cat}(\alpha_t y + (1 - \alpha_t) m). \quad (25)$$

Decoding proceeds from $t=1$ to $t=0$ in T steps. Let $\pi_\theta(y_t) = \text{softmax}(f_\theta(y_t))$ and $0 \leq s < t \leq 1$ with $s, t \in \{0, \frac{1}{T}, \dots, 1\}$. Define

$$\lambda_{s|t} = \frac{\alpha_s - \alpha_t}{1 - \alpha_t} \in [0, 1]. \quad (26)$$

Then each position i evolves by the copy/denoise rule

$$\text{copy: } y_t^{(i)} \neq M \Rightarrow y_s^{(i)} = y_t^{(i)}, \quad \text{denoise: } y_t^{(i)} = M \Rightarrow y_s^{(i)} \sim (1 - \lambda_{s|t}) m + \lambda_{s|t} \pi_\theta(y_t). \quad (27)$$

In practice we perform parallel position updates. Optional heuristics (e.g., blockwise semi-autoregressive ordering, confidence-based remasking) can be added without changing Eqs. equation 25–equation 27.

E DETAILED EXPERIMENTS SETTINGS

Our evaluation is conducted using a framework adapted from the `lm-evaluation-harness`, tailored for diffusion language models. As mentioned in the main text, our benchmarks are grouped into three categories: general-knowledge QA, mathematical and scientific reasoning, and code generation. The specific settings for each model family are detailed below. For all code generation tasks, the `HF_ALLOW_CODE_EVAL=1` environment variable is set to enable execution-based evaluation.

E.1 SETTINGS FOR LLaDA MODELS

For the LLaDA model series (LLaDA-8B-Base, LLaDA-8B-Instruct, and LLaDA-1.5), we apply a consistent set of hyperparameters for each evaluation task to ensure a fair comparison.

General-knowledge QA. For tasks based on conditional likelihood estimation, we use 128 Monte Carlo samples ($mc_{num} = 128$) unless specified otherwise. The task-specific settings are as follows:

- **MMLU:** We use a 5-shot setting with a Classifier-Free Guidance (CFG) scale of 0.0. For our ablation runs, mc_{num} was set to 1 for faster evaluation.
- **WinoGrande:** We use a 5-shot setting with a CFG scale of 0.0.
- **HellaSwag, ARC-Easy, ARC-Challenge:** These are evaluated in a 0-shot setting with a CFG scale of 0.5.
- **PIQA:** This is evaluated in a 0-shot setting with a CFG scale of 0.5.

Mathematical & Scientific Reasoning and Code Generation. For tasks requiring conditional generation, the settings are configured to balance performance and generation quality.

- **GSM8K:** We employ diffusion semi-autoregressive sampling by setting $gen_{length} = 256$ and a $block_{length} = 32$, with 256 denoising steps. This approach is effective for complex, multi-step reasoning.
- **BBH:** This task is evaluated with standard parallel decoding, where gen_{length} , $steps$, and $block_{length}$ are all set to 128.
- **GPQA:** This task is evaluated as a likelihood estimation problem in a 0-shot setting, with a batch size of 1, $mc_{num} = 128$, and a CFG scale of 0.5.
- **HumanEval & MBPP:** Both code generation tasks are evaluated with standard parallel decoding, where gen_{length} , $steps$, and $block_{length}$ are all set to 128.

E.2 SETTINGS FOR DREAM MODELS

For the Dream model series (Dream-7B-Base and Dream-7B-Instruct), a similar evaluation protocol is followed to maintain consistency. The primary difference is in the number of few-shot examples for some tasks, as detailed in the evaluation script.

General-knowledge QA. Like the LLaDA series, these tasks are evaluated using 128 Monte Carlo samples. The few-shot settings are as follows:

- **MMLU, WinoGrande:** Evaluated in a 5-shot setting.
- **ARC-Easy, ARC-Challenge, HellaSwag, PIQA:** All evaluated in a 0-shot setting.

Table 1: Effectiveness of MCS on previous methods. We report MMLU in 5 shots.

((a)) LLaDA-8B-Base				((b)) Dream-7B-Base			
Model	MCS	Method	MMLU(5) ↑	Model	MCS	Method	MMLU(5) ↑
LLaDA-8B-Base	✗	GPTQ	34.75	Dream-7B-Base	✗	GPTQ	23.84
	✓	GPTQ	36.75		✓	GPTQ	24.08
	✗	GPTAQ	34.73		✗	GPTAQ	23.90
	✓	GPTAQ	37.92		✓	GPTAQ	24.68
	✗	Slim-LLM	47.98		✗	Slim-LLM	25.54
	✓	Slim-LLM	50.11		✓	Slim-LLM	26.01

Mathematical & Scientific Reasoning and Code Generation. For generative tasks, the Dream models are evaluated using the same general protocol as LLaDA.

F MORE EXPERIMENTS

F.1 EFFECTIVENESS OF MCS ON BASELINE METHODS

To demonstrate the general applicability of our Masked Calibration Simulation (MCS), we investigate its effectiveness when applied to existing state-of-the-art (SOTA) post-training quantization methods. We integrate MCS into the calibration stage of GPTQ, GPTAQ, and Slim-LLM, and evaluate the impact on MMLU 5-shot accuracy for the LLaDA-8B-Base and Dream-7B-Base models. The results are presented in Table 1.

As shown in Table 1(a), applying MCS consistently improves the performance of all baseline methods on LLaDA-8B-Base. Specifically, the MMLU score for GPTQ increases from 34.75% to 36.75%, and GPTAQ sees a significant boost from 34.73% to 37.92%. Similarly, the performance of the strong baseline, Slim-LLM, is enhanced from 47.98% to 50.11%.

A similar trend is observed on the Dream-7B-Base model, as detailed in Table 1(b). While the improvements are more modest on this model, MCS still provides consistent gains across the board. For instance, GPTAQ’s accuracy rises from 23.90% to 24.68%, and Slim-LLM’s score improves from 25.54% to 26.01%.

These results strongly indicate that MCS is not limited to our DAQ framework but serves as a valuable, standalone module for quantizing dLLMs. By aligning the calibration data distribution with the model’s native diffusion-denoising process, MCS effectively reduces the calibration-inference mismatch and boosts the performance of various PTQ techniques.