

EFFICIENT BLOCKWISE DIVERSE ACTIVE LEARNING

Anonymous authors

Paper under double-blind review

ABSTRACT

Active learning (AL) techniques are known for selecting the most informative data points from large datasets, thereby enhancing model performance with fewer labeled samples. This makes AL particularly useful in tasks where labeling is limited or resource-intensive. However, most existing effective methods rely on uncertainty scores to select samples, often overlooking diversity, which results in redundant selections, especially when the batch size is small compared to the overall dataset. This paper introduces Efficient Blockwise Diverse Active Learning (EBDAL), a generalizable framework that combines uncertainty with diversity-based selection to overcome these limitations. By partitioning the dataset into blocks via a clustering strategy, we ensure diverse sampling within each block, enabling more efficient handling of large-scale datasets. To quantify diversity, we minimize the Maximum Mean Discrepancy (MMD) between the selected subset and the full dataset, which is then reformulated as a Quadratic Unconstrained Binary Optimization (QUBO) problem. The resulting QUBO is submodular, which permits an efficient greedy algorithm. We further demonstrate feasibility on real quantum hardware through an end-to-end selection experiment. Our experimental results demonstrate that EBDAL not only improves the accuracy of uncertainty-based strategies but also outperforms a wide range of selection methods, achieving substantial computational speedups. The findings highlight EBDAL’s robustness, efficiency, and adaptability across various datasets.

1 INTRODUCTION

Active learning (AL) is a learning paradigm in which a model interactively requests labels for the most informative samples from a large unlabeled pool, aiming to attain strong predictive performance with substantially fewer annotations. This is particularly relevant in modern deep learning, where unlabeled data are plentiful but labeling is costly, time-consuming, or requires domain expertise (e.g., medical imaging, autonomous driving). By selectively querying labels while maintaining accuracy, AL provides a practical route to scaling supervised learning under budget constraints. Methodologically, AL spans uncertainty sampling and Bayesian approximations (Lewis & Gale, 1994; Gal et al., 2017; Yoo & Kweon, 2019), geometric/coreset and representative selection (Sener & Savarese, 2018), gradient-based hybrid criteria (Ash et al., 2019), and variational/adversarial formulations (Sinha et al., 2019), with comprehensive surveys available in (Settles, 2009; Li et al., 2024). Collectively, these lines of work underscore the need for query strategies that jointly balance uncertainty, representativeness, and computational scalability.

Uncertainty-based active learning has achieved strong empirical success, reliably reducing annotation cost by prioritizing highly informative queries. Nevertheless, redundancy frequently arises under tight budgets: top-uncertainty samples often cluster within the same local region, leading to inefficient label usage (Figure 1). Recent efforts have begun to address these issues but leave important gaps. For example, Wang et al. (2024b) automates strategy choice from a set of heuristics, advancing ease of use; yet committing to a single policy per round limits the ability to balance uncertainty and diversity within a batch, and the strategy-search/preparation phases introduce nontrivial overhead that is often omitted from end-to-end timing. Likewise, Bae et al. (2025) improves representativeness via uncertainty-weighted coverage, but the approach remains sensitive to uncertainty calibration, is still prone to redundancy when uncertainty mass concentrates, and can under-sample globally representative regions with lower uncertainty. These observations suggest that a more effective solution should explicitly integrate uncertainty with distributional diversity while remaining computationally efficient across datasets and budget regimes.

To address these challenges, we propose Efficient Blockwise Diverse Active Learning (EBDAL), a scalable framework that couples uncertainty with distributional diversity. We operationalize diversity through maximum mean discrepancy (MMD), selecting a subset that approximates the kernel mean of the unlabeled pool, thereby promoting broad coverage and reducing redundancy. Instead of performing a single global optimization, EBDAL adopts a blockwise MMD minimization strategy, significantly improving time and memory efficiency. We establish a bounded approximation gap between the blockwise objective and the global optimum. To balance informativeness and representativeness, our method selects samples from a high-uncertainty candidate set without over-relying on the precise uncertainty values. Technically, the blockwise MMD selection is formulated as a Quadratic Unconstrained Binary Optimization (QUBO) problem with a submodular objective. This formulation admits an efficient greedy algorithm in practice and is also compatible with modern quantum hardware.

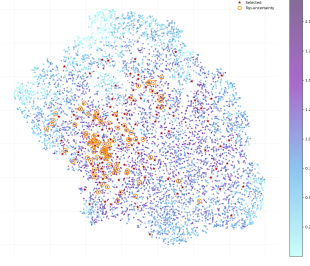


Figure 1: Visualization example of redundancy

The main contributions in this work are as follows:

- **A scalable, redundancy-aware AL framework.** We introduce **EBDAL**, which integrates uncertainty with *distributional diversity* through *blockwise* MMD minimization. Operating per block promotes broad feature-space coverage and substantially reduces small-batch redundancy common to uncertainty-only selection, while remaining simple to deploy on large unlabeled pools.
- **Theory and efficiency.** We provide a *bounded approximation guarantee* showing that blockwise MMD closely tracks the global MMD objective, preserving distributional alignment. In practice, our complexity analysis and wall-clock experiments show consistent speedups compared with full-pool, without loss of accuracy.
- **Quantum-enabled selection.** We formulate the blockwise MMD selection as a QUBO problem whose objective is submodular. This formulation enables two complementary solving approaches: (i) a fast greedy algorithm for practical applications, and (ii) direct access to real quantum hardware. To our knowledge, this yields the first demonstration of an end-to-end active-learning loop that uses a quantum computer for the selection phase and a classical computer for model training, applied to classical datasets.

2 RELATED WORK

Active Learning Active learning (AL) targets high performance under tight labeling budgets by querying the most informative points. Early work formalized uncertainty sampling for text (Lewis & Gale, 1994). With deep models, approximate Bayesian inference enabled mutual-information criteria such as Bayesian Active Learning by Disagreement (Gal et al., 2017); learnable acquisition functions predict per-sample loss (Yoo & Kweon, 2019); and batch-diverse methods curb redundancy via gradient-space clustering (BADGE) (Ash et al., 2019) or variational adversarial latent-space selection (Sinha et al., 2019). Deterministic uncertainty estimation (DDU) provides a strong non-Bayesian baseline (Mukhoti et al., 2023). Recent efforts emphasize scalability and automation, including leverage-score sampling (Shimizu et al., 2023), uncertainty herding (Bae et al., 2025), and differentiable strategy search (Wang et al., 2024b).

Coreset Selection Coresets approximate large datasets with small, weighted summaries to accelerate training while maintaining accuracy. In AL, geometry-driven k -center selection reduces redundancy for CNNs (Sener & Savarese, 2018). More broadly, coreset techniques offer scalable summaries with provable guarantees (Mirzasoleiman et al., 2020) and bilevel updates for streaming and continual learning (Borsos et al., 2020). Recent diversity-centric subset design revisits coverage for robustness (e.g., instruction tuning) (Wang et al., 2024a). Our approach is coreset-like in spirit but differs by optimizing an RKHS-grounded diversity objective (MMD) in a blockwise, active pipeline that includes uncertainty filtering and an acceleration path via a QUBO reformulation.

Submodularity and QUBO/Ising For subset selection under a cardinality constraint, constructing (approximately) submodular objectives is a powerful design principle. Scalable variants tailor this idea to machine learning, e.g., Bayesian batch AL as sparse subset approximation (Pinsler et al., 2019), and distributed optimization for pairwise submodular functions (Böther et al., 2024). A complementary line formulates selection as QUBO/Ising (see modeling and mapping references (Lucas, 2014; Kochenberger et al., 2014; Glover et al., 2018)), enabling specialized solvers and quantum optimizers such as quantum annealing (Kadowaki & Nishimori, 1998) and QAOA (Farhi et al., 2014). In contrast to (Bae et al., 2025), which relies on uncertainty-weighted coverage and can be sensitive to miscalibration and cluster concentration, our method optimizes an explicit MMD-based distributional alignment with blockwise structure and a provable approximation to the global objective, yielding stronger redundancy mitigation and robustness.

3 METHODOLOGY

3.1 PROBLEM DEFINITION

We consider a standard pool-based active learning setup. Let $\mathcal{D}_U = \{\mathbf{x}_i\}_{i=1}^N$ denote a large pool of unlabeled data instances. The learning process begins with a small labeled set $\mathcal{D}_L = \{(\mathbf{x}_j, y_j)\}_{j=1}^M$, where $M \ll N$ and y_j represents the ground-truth label for $\mathbf{x}_j$. The core of the problem is a sequential querying procedure governed by a fixed annotation budget Q . In each iteration t , an acquisition function $\alpha(\mathbf{x}, \mathcal{M}_{t-1})$, which quantifies the utility of querying the label of a point $\mathbf{x}$, is used to select a batch of b samples from the unlabeled pool:

$$\mathcal{D}_q^t = \underset{\mathcal{D} \subset \mathcal{D}_U, |\mathcal{D}|=b}{\operatorname{argmax}} \sum_{\mathbf{x} \in \mathcal{D}} \alpha(\mathbf{x}, \mathcal{M}_{t-1}), \quad (1)$$

The oracle then provides the labels $\{y\}$ for the queried batch $\mathcal{D}_q^t$. These newly labeled samples are removed from $\mathcal{D}_U$ and added to $\mathcal{D}_L$. The model $\mathcal{M}_{t-1}$ is then trained on the newly labeled batch $\mathcal{D}_q^t$, and the training cycle repeats. This select-query-retrain cycle continues until the annotation budget Q is exhausted. The objective is to maximize the sample efficiency of the learning process, obtaining a highly accurate model with as few labeled examples as possible.

To mitigate the redundancy of conventional global top- b uncertainty sampling, we adopt a filter-then-match paradigm that couples uncertainty with diversity. We first retain a generous set of high-uncertainty candidates, yielding a filtered search pool

$$\mathcal{D}_U^{\text{flt},t} = \{x \in \mathcal{D}_U \mid x \text{ is among the } k \text{ most uncertain points at iteration } t\}, \quad k \gg b. \quad (2)$$

This step uses uncertainty purely as a screen: it prunes clearly low-utility points and curbs local redundancy, while preserving a high-recall set of plausible candidates. Given the screened pool, the final batch is selected by aligning with the distribution of the entire unlabeled set via MMD:

$$\mathcal{D}_q^t = \underset{\mathcal{D} \subset \mathcal{D}_U^{\text{flt},t}, |\mathcal{D}|=b}{\operatorname{argmin}} \operatorname{MMD}^2(\mathcal{D}, \mathcal{D}_U). \quad (3)$$

Uncertainty forms a high-recall candidate set, and MMD enforces coverage/representativeness in the final picks—yielding batches that are both informative and diverse. To scale to large datasets, we first cluster the pool into feature-space blocks and then apply Eq.(2) and Eq.(3) in parallel within each block before aggregating the selections.

3.2 DIVERSITY-ENHANCED

Active learning under a fixed budget can be viewed as iterative coresets selection: at each round, we choose a small subset that minimizes the expected training loss (i.e., the loss gap relative to using the full pool). This perspective turns batch acquisition into selecting representative points that preserve loss, naturally motivating the joint use of uncertainty and coverage/diversity. Formally, the objective can be expressed as:

$$\min_{\mathcal{D}_q^t: |\mathcal{D}_q^t|=b} \mathbb{E}_{x,y \sim P} [\ell(x, y; \mathcal{M}_t)], \quad (4)$$

where $\mathbf{x}_i, y_i$ are i.i.d. drawn from an underlying distribution P , and $l(x, y; \mathcal{M}_{t-1})$ is the loss function of the model $\mathcal{M}_{t-1}$ after being trained on the batch $\mathcal{D}_q^t$.

Directly minimizing the training loss in Eq. (4) is computationally challenging. To overcome this, we introduce a tractable surrogate objective: instead of minimizing the loss directly, we propose to minimize the MMD between the selected subset and the full dataset. We establish a connection between the upper and lower bounds of the training loss function over a given subset and the MMD of that subset with respect to the full unlabeled pool. Specifically, it measures the distance between the distribution of the selected coreset and that of the full unlabeled set. The formal statement of MMD between two distributions P and Q is as follows:

$$\text{MMD}[\mathcal{F}, P, Q] = \sup_{\|f\|_{\mathcal{H}} \leq 1} (\mathbb{E}_P[f(x)] - \mathbb{E}_Q[f(y)]), \quad (5)$$

where $\mathcal{F}$ is a class of functions, and $\|f\|_{\mathcal{H}}$ denotes the norm of the function f in the reproducing kernel Hilbert space (RKHS). The expectations $\mathbb{E}_P[f(x)]$ and $\mathbb{E}_Q[f(y)]$ represent the expected values of the function f under the distributions P and Q . The MMD measures the largest difference between the distributions P and Q in terms of the function class $\mathcal{F}$. By minimizing the MMD, we encourage the selected subset to closely match the full distribution of the unlabeled data, ensuring that the selected coreset is a good approximation of the entire dataset. Before formally establishing the connection between the MMD discrepancy and the expected loss, we first state the following key lemma.

Lemma 1. (Zheng et al., 2022) *Given n i.i.d. samples drawn from P_μ as $S = \{(x_i, y_i)\}_{i \in [n]}$, where $y_i \in [C]$ is the class label for example x_i , a coreset S' that is a p -partial r -cover for P_μ on the input space X , and any $\varepsilon > 1 - p$, suppose: (i) the loss $l(\cdot, y, w)$ is λ_ℓ -Lipschitz continuous for all y, w and bounded by L ; (ii) the class-specific regression function $\eta_c(x) := \mathbb{P}(y = c \mid x)$ is λ_η -Lipschitz for all $c \in [C]$; and (iii) $l(x, y; h_{S'}) = 0$ for all $(x, y) \in S'$. Then, with probability at least $1 - \varepsilon$,*

$$\left| \frac{1}{n} \sum_{(x,y) \in S} l(x, y; h_{S'}) \right| \leq r(\lambda_\ell + \lambda_\eta LC) + L \sqrt{\frac{\log\left(\frac{p}{p+\varepsilon-1}\right)}{2n}}. \quad (6)$$

Informally, we interpret a p -partial r -cover as the fraction p of points in the ground set X that fall within distance r of the subset S (the formal definition is deferred to the Appendix A). By the preceding lemma, with probability at least p , the loss in a single training pass is upper-bounded by a quantity that is monotone in the covering radius r , and therefore is primarily determined by r . Hence minimizing the loss can be viewed as selecting points with strong covering power under the p -partial r -cover criterion. To facilitate such selection, we relate coverage to the empirical MMD below.

Theorem 1 (Coverage–MMD relation). *Let X be the full dataset and $S \subseteq X$ a subset. For a radius $r > 0$, denote the coverage ratio by*

$$\rho(r) := \frac{1}{|X|} |\{x \in X : \text{dist}(x, S) \leq r\}|.$$

Let $k(\cdot, \cdot)$ be the Gaussian kernel $k(u, v) = \exp(-\|u - v\|^2 / (2\sigma^2))$. Then the following bounds hold:

$$\frac{\delta_{\text{MMD}} - e^{-r^2/(2\sigma^2)}}{1 - e^{-r^2/(2\sigma^2)}} < \rho(r) < \delta_{\text{MMD}} |S| e^{r^2/(2\sigma^2)}. \quad (7)$$

Here δ_{MMD} denotes the cross term in the empirical MMD between S and X .

Theorem 1 establishes a precise connection between the MMD term δ_{MMD} and data coverage. The bounds in Eq.(7) show that a larger δ_{MMD} (implying a smaller MMD) tightens both the lower and upper bounds on the coverage ratio $\rho(r)$ for any radius r . This means that for a fixed target coverage level p , the required radius r_p decreases as δ_{MMD} increases. Combining this with Lemma 1 yields the key insight: minimizing the MMD leads to a tighter probabilistic loss bound for a given success probability p , by reducing the effective covering radius needed.

3.3 FROM MMD-BASED DIVERSITY TO A QUBO FORMULATION

When P and Q are represented by empirical samples $X = \{x_i\}_{i=1}^n$ and $S \subseteq X$, and using the kernel mean embeddings $\mu_X = \frac{1}{n} \sum_{i=1}^n \phi(x_i)$ and $\mu_S = \frac{1}{|S|} \sum_{x \in S} \phi(x)$, the squared MMD admits a closed-form expression:

$$f(S) = \|\mu_S - \mu_X\|_{\mathcal{H}}^2.$$

Let $K \in \mathbb{R}^{n \times n}$ be the Gram matrix with $K_{ij} = \phi(x_i) \cdot \phi(x_j) = k(x_i, x_j)$ and $m \in \{0, 1\}^n$ encode the selected subset with $\mathbf{1}^\top m = k$. The MMD-minimization problem is

$$\min_{m \in \{0, 1\}^n} \frac{1}{k^2} m^\top K m - \frac{2}{nk} \mathbf{1}^\top K m \quad \text{s.t.} \quad \mathbf{1}^\top m = k. \quad (8)$$

With a sufficiently large penalty on the cardinality constraint, Eq.(8) is directly converted into a QUBO of the form $\min_{m \in \{0, 1\}^n} m^\top \hat{Q} m$ having the same minimizers; the construction of $\hat{Q}$ and the equivalence proof are deferred to Appendix B. Moreover, since common kernels satisfy $K_{ij} \geq 0$, maximizing the negated objective of Eq.(8) under the same cardinality constraint yields a quadratic set function whose continuous relaxation has the Hessian with non-positive off-diagonal entries, hence it is submodular (Bian et al., 2017). Therefore, the instance can be tackled either by (cardinality-constrained) greedy selection or by quantum QUBO/Ising solvers (see Appendix B).

Even with extensive engineering, greedy selection over the full pool is costly. If we precompute the non-sparse Gram matrix once and update marginal gains incrementally, the computational complexity of k -step greedy procedure on n points is as follows:

$$T_{\text{greedy}}^{\text{full}} = O(n^2 + kn). \quad (9)$$

After partitioning $X = \bigcup_{b=1}^B X_b$ with $|X_b| = n_b$ and running greedy *inside each block*, the total cost becomes

$$T_{\text{greedy}}^{\text{block}} = O\left(\sum_{b=1}^B n_b^2 + \sum_{b=1}^B k_b n_b\right), \quad \sum_{b=1}^B k_b = k. \quad (10)$$

When blocks are reasonably balanced ($n_b \approx n/B$ and $k_b \approx k/B$),

$$T_{\text{greedy}}^{\text{block}} = O\left(\frac{n^2}{B} + \frac{kn}{B}\right), \quad (11)$$

the QUBO in Eq.(8) involves one binary variable per data point, making it prohibitively slow and yielding poor solutions on large datasets for both classical and quantum solvers. Our blockwise partitioning confines each QUBO instance to a scale of a few hundred variables, which is well within the regime where classical optimizers perform efficiently and effectively, and is also compatible with the capabilities of current quantum hardware.

To strengthen the alignment between a small MMD and strong coverage, as suggested by Theorem 1 which becomes tighter for data blocks with smaller diameter, we partition the data into compact clusters (e.g., in kernel feature space) using K-means for its simplicity and scalability. Moreover, this blockwise approach is principled: when the selection quota for each block is set proportional to its size, optimizing MMD separately within each block serves as a surrogate for the global objective, since the latter is governed by a weighted average of the per-block objectives (a generalized statement with a quota-mismatch penalty is given in Appendix C).

Theorem 2 (Blockwise MMD approximation with proportional quotas). *Let $X = \bigcup_{b=1}^B X_b$ be a partition with weights $w_b := |X_b|/|X|$. From each block b , select $S_b \subseteq X_b$ of size $|S_b| = k_b$ and set $S = \bigcup_{b=1}^B S_b$, with $\alpha_b := k_b/k$ and $\sum_{b=1}^B \alpha_b = \sum_{b=1}^B w_b = 1$. Let k be a bounded positive definite kernel with feature map ϕ into an RKHS $\mathcal{H}$, and assume $\sup_x k(x, x) \leq \kappa^2$ (for a normalized RBF kernel, $\kappa = 1$). Denote $\mu_A := |A|^{-1} \sum_{x \in A} \phi(x)$ and $\text{MMD}^2(A, B) := \|\mu_A - \mu_B\|_{\mathcal{H}}^2$. If quotas are proportional, i.e., $\alpha = w$, then*

$$\text{MMD}^2(S, X) \leq 2 \sum_{b=1}^B w_b \text{MMD}^2(S_b, X_b) \leq 2 \max_b \text{MMD}^2(S_b, X_b). \quad (12)$$

In particular, if each blockwise procedure achieves $\text{MMD}^2(S_b, X_b) \leq \varepsilon_b$, then $\text{MMD}^2(S, X) \leq 2 \sum_b w_b \varepsilon_b \leq 2 \max_b \varepsilon_b$.

Theorem 2 indicates that with proportional quota allocation ($\alpha = w$), the global MMD error is bounded by a weighted average of the per-block errors. Thus, by ensuring the selection from each compact cluster is representative, the blockwise procedure closely approximates the global objective. Consequently, optimizing MMD within each block is a principled surrogate for the intractable global optimization. If the quotas deviate from this proportion, an additional penalty term is incurred.

3.4 ALGORITHM

Given a candidate set $C_b \subseteq X_b$ and quota k_b , we solve the per-block MMD problem

$$\min_{S_b \subseteq C_b, |S_b|=k_b} \text{MMD}^2(S_b, X_b) = \left\| \frac{1}{k_b} \sum_{x \in S_b} \phi(x) - \frac{1}{n_b} \sum_{x \in X_b} \phi(x) \right\|_{\mathcal{H}}^2. \quad (13)$$

With proportional quotas ($k_b \propto |X_b|$), Theorem 2 ensures that the global error is controlled by a weighted average of the per-block errors.

Algorithm 1 EBDAL (multi-round active learning)

Require: Initial labeled set L_0 , unlabeled pool U_0 , number of AL rounds T , per-round batch size k , candidate ratio τ , max block size $M_{\max}$, kernel $k(\cdot, \cdot)$, uncertainty scorer $u_\theta(\cdot)$, SOLVER for Eq. equation 13

Ensure: Final trained model θ_T

```

1:  $L \leftarrow L_0, U \leftarrow U_0$ 
2: for  $t = 1$  to  $T$  do ▷ outer AL loop
3:   Train / fine-tune model parameters  $\theta$  on  $L$ 
4:   Compute embeddings  $z_\theta(x)$  for all  $x \in U$ 
5:   Cluster  $U$  in the kernel/feature space (e.g., recursive Kmeans) so each block size  $\leq M_{\max}$ :
      $U = \bigcup_{b=1}^B X_b, n_b \leftarrow |X_b|$ 
6:   Set block quotas:  $k_b \leftarrow \text{round}(k n_b / |U|)$ ; adjust by  $\pm 1$  to ensure  $\sum_b k_b = k$ 
7:   Set global candidate budget  $M \leftarrow \lceil \tau |U| \rceil$  and distribute to blocks:  $c_b \leftarrow \min\{n_b, \max(k_b, \text{round}(M n_b / |U|))\}$ ; adjust by  $\pm 1$  so  $\sum_b c_b = M$ 
8:    $S^{(t)} \leftarrow \emptyset$ 
9:   for  $b = 1$  to  $B$  do
10:    Rank  $X_b$  by uncertainty  $u_\theta(\cdot)$  (descending); let  $C_b$  be the top- $c_b$  in  $X_b$ 
11:    (Optional) Fit per-block kernel parameters (e.g., RBF bandwidth via the median heuristic on  $X_b$ )
12:    Solve the per-block objective:  $S_b \leftarrow \text{SOLVER}(C_b, X_b, k_b; k(\cdot, \cdot))$  for Eq. (13)
13:     $S^{(t)} \leftarrow S^{(t)} \cup S_b$ 
14:    Query labels  $y$  for all  $x \in S^{(t)}$  from the oracle
15:     $L \leftarrow L \cup \{(x, y) : x \in S^{(t)}\}, U \leftarrow U \setminus S^{(t)}$ 
16:    Train model  $f_{\theta_T}$  using  $\{(x, y) : x \in S^{(t)}\}$ 
17: return  $\theta_T$ 

```

Remarks. (i) By setting quotas proportional to block sizes ($k_b \propto n_b$), we adhere to the conditions of Theorem 2, which eliminates the quota-mismatch penalty and bounds the global MMD by a weighted sum of the local block MMDs.

(ii) Candidate filtering ($c_b \ll n_b$) reduces the kernel computation cost within each block from $O(n_b^2)$ to $O(c_b^2)$ and significantly shrinks the problem size for the downstream solver.

(iii) Blocks are processed independently, enabling parallelization. The framework is solver-agnostic, accommodating various backends (e.g., QUBO/Ising solvers; see Appendix D).

(iv) **Hardware-aware Ising mapping.** To cope with the limited bitwidth and dynamic range of current quantum hardware, we adopt a precision-aware multi-auxiliary construction that *splits large data-spin fields* into several couplers to a small set of ancilla spins. This redistributes scale so that individual coefficients are comparable to the typical J_{ij} level after global rescaling/quantization, thereby preserving the effective resolution of J and mitigating precision loss; implementation choices (e.g., $\Delta_J, \gamma, G_i, B_g$) are detailed in Appendix D.

4 EXPERIMENTS

We evaluate our diversity-enhanced acquisition framework on standard image-classification benchmarks by integrating it with representative uncertainty criteria and comparing each baseline to its diversity-enhanced counterpart, reporting both accuracy-budget curves and the Area Under the Curve (AUC). We also benchmark the selection backend by comparing a fast greedy solver against a quantum QUBO/Ising formulation. Finally, we perform ablations that successively enable three core modules: block partitioning with proportional per-block budgets; an in-block uncertainty filter that forms a candidate pool; and MMD-based selection within each block. This setup isolates their individual contributions.

4.1 EXPERIMENT SETTINGS

Datasets and Baselines We evaluate on five image classification benchmarks: CIFAR-10, CIFAR-100 (Krizhevsky et al., 2009), MNIST (Deng, 2012), SVHN (Netzer et al., 2011), and Fashion-MNIST (Xiao et al., 2017). To study the benefit of diversity, we combine our framework with three representative uncertainty-based selectors: *entropy* (probability-based) (Shannon, 2001), *BALD* (Bayesian active learning by disagreement) (Gal et al., 2017), and *LPL* (learning-loss predictor) (Yoo & Kweon, 2019). Methods suffixed with “-EBDAL” denote the corresponding algorithms augmented by our diversity module.

As baselines, we include the standalone versions of the above three uncertainty methods (*entropy*, *BALD*, *LPL*), as well as *BADGE* (Ash et al., 2019), *Kmeans* (Lloyd, 1982) clustering-based selection, and the recent uncertainty-herding approach *Uherding* (Bae et al., 2025). Further implementation details (initial labeled size, labeling budgets, model/backbone, training schedules, and hyperparameters) are provided in the Appendix E.

4.2 PERFORMANCE COMPARISON ACROSS DATASETS

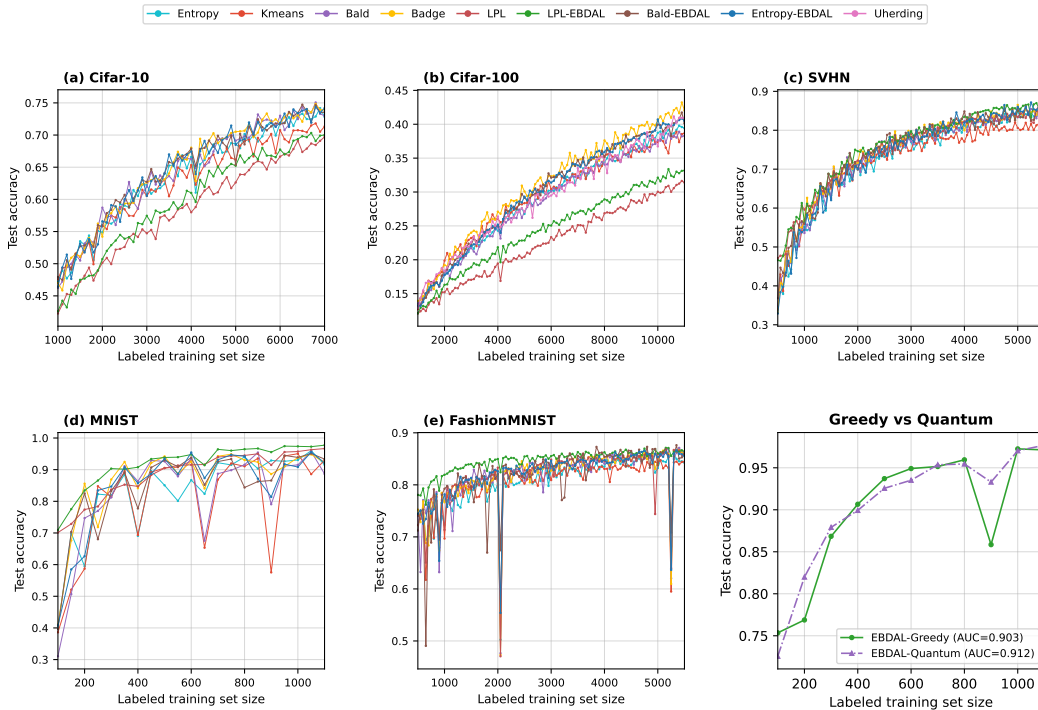


Figure 2: Overall performance trends across labeling budgets. Test accuracy vs. labeled training set size on five benchmarks.

Method	CIFAR-10	CIFAR-100	SVHN	MNIST	FMNIST
Bald	0.6450	0.2866	0.7351	0.8392	0.8110
Bald-EBDAL	0.6482 (+0.32)	0.3005 (+1.39)	0.7440 (+0.89)	0.8589 (+1.97)	0.8180 (+0.70)
Entropy	0.6391	0.2902	0.7259	0.8418	0.8061
Entropy-EBDAL	0.6469 (+0.78)	0.2997 (+0.95)	0.7359 (+1.00)	0.8573 (+1.55)	0.8212 (+1.51)
LPL	0.5851	0.2258	0.7478	0.8880	0.8184
LPL-EBDAL	0.5977 (+1.26)	0.2446 (+1.88)	0.7595 (+1.17)	0.9219 (+3.39)	0.8467 (+2.83)
Badge	0.6480	0.3107	0.7448	0.8748	0.8188
Kmeans	0.6282	0.2940	0.7209	0.8146	0.8039
Uherding	0.6441	0.2923	/	/	/

Table 1: AUC (Area Under Curve) across datasets. Best per dataset in **bold**. Values in parentheses show the improvement over the corresponding non-EBDAL baseline (in percentage points).

Across the five image classification benchmarks, we observe a consistent pattern: for a fixed labeling budget, uncertainty methods augmented with our diversity module (the “-EBDAL” variants) shift the entire accuracy–budget curve upward relative to their uncertainty-only counterparts. This translates into fewer redundant queries and more effective use of labels (see Figure 2; quantitative AUCs appear in Table 1).

The gains are especially pronounced on settings with many classes and noisier raw uncertainty signals (e.g., CIFAR-100), where both *LPL-EBDAL* and *Entropy-EBDAL* deliver clear AUC improvements over their baselines. On easier datasets such as MNIST and SVHN, where most methods already perform strongly at small budgets, EBDAL variants still yield consistent positive shifts; notably, *LPL-EBDAL* attains the best AUC on MNIST.

Compared with diversity-only baselines (e.g., *KMeans*, *BADGE*), EBDAL achieves a more robust exploration–exploitation balance by first forming a high-uncertainty candidate pool within each block and then selecting a subset that minimizes MMD to the block population. This design avoids small-batch redundancy caused by concentrated uncertainty mass while preserving global representativeness and scalability.

The sixth panel further contrasts greedy and quantum backends under the same setup. The two curves closely align, with AUCs of 0.903 (*EBDAL-Greedy*) and 0.912 (*EBDAL-Quantum*), demonstrating the feasibility of an end-to-end selection–training loop on real hardware. In our testbed, the quantum backend solves a 500-variable QUBO instance in about 0.2 ms per call, which is orders of magnitude faster than the second-scale runtime of the greedy (submodular) solver on problems of comparable size. Given the current invocation associated with quantum resources, we therefore use the greedy solver for the main experiments and ablations, and report the quantum results as a feasibility and effectiveness reference. A representative Hamiltonian-energy evolution trace and brief explanation are provided at the end of Appendix E.

Overall, EBDAL delivers consistent accuracy gains and strong computational scalability across datasets and budgets, validating the benefit of integrating blockwise diversity with uncertainty in practical active learning pipelines.

4.3 ABLATION STUDIES

We ablate the three components of our diversity-enhanced acquisition pipeline—(i) block partitioning in feature space, (ii) per-block uncertainty filtering that forms an uncertainty candidate pool, and (iii) distributional selection via in-block MMD minimization—using three variants: *LPL (standard)* without blocks or explicit diversity; *Block-TopU* with block partitioning and UCPs but selecting the k_b most-uncertain points per block; and *LPL-EBDAL* with the full blockwise MMD selection. All configurations share the same backbone, training schedule, data splits, and label budgets; curves report test accuracy versus labeled-set size to isolate acquisition effects.

Across MNIST, CIFAR-10, and SVHN, the full *LPL-EBDAL* consistently dominates its ablations over a wide range of budgets, indicating that explicit distributional matching (MMD) adds value beyond uncertainty guidance and block partitioning alone. *Block-TopU* improves over *LPL* in the low-budget regime by dispersing queries across blocks (reducing obvious redundancy), but it tends

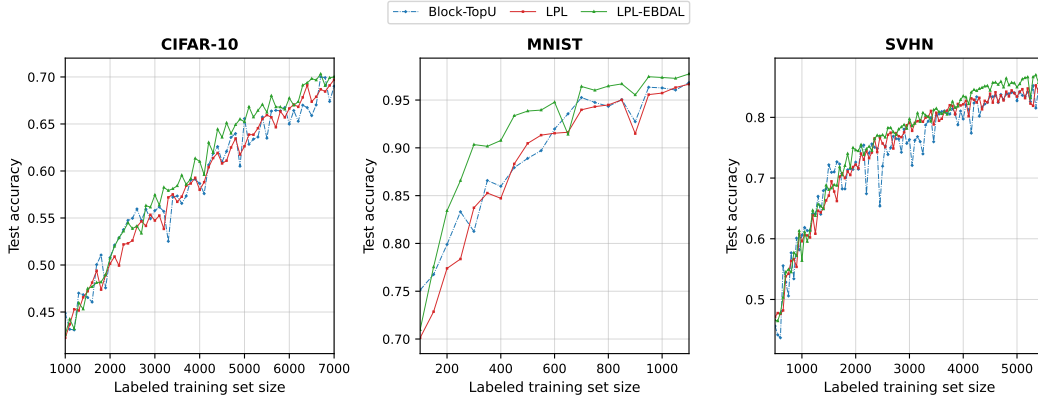


Figure 3: Ablation on *MNIST*, *CIFAR10*, and *SVHN*. We compare **LPL** (no blocking/diversity), **Block-TopU** (blocking + per-block top-uncertainty), and **LPL-EBDAL** (full). **LPL-EBDAL** consistently achieves the best label-efficiency by combining uncertainty guidance with MMD-based distributional matching within each block.

to plateau earlier when high-uncertainty points cluster within the same local regions. In contrast, *LPL-EBDAL* continues to benefit as the budget grows, suggesting that aligning the selected subset’s kernel mean with each block’s population mitigates within-block redundancy and recovers globally representative samples that pure uncertainty may undervalue. Representative trajectories are shown in Figure 3.

5 CONCLUSION

We presented Efficient Blockwise Diverse Active Learning (EBDAL), a practical framework that integrates uncertainty guidance with distributional diversity through blockwise MMD minimization. By framing active learning as iterative coreset selection, we establish a theoretical link between distributional alignment and loss control. Technically, we formulate the per block MMD objective as a submodular QUBO, enabling efficient greedy selection and compatibility with quantum solvers. Theoretically, we derive a coverage MMD bound and prove that blockwise optimization with proportional quotas closely approximates the global objective. Empirically, EBDAL consistently improves accuracy budget curves and AUC over uncertainty only and diversity only baselines across standard image benchmarks, while reducing runtime. Ablations validate the contribution of each component, and a direct comparison demonstrates the feasibility of executing the selection phase on real quantum hardware.

In summary, EBDAL delivers a scalable and label efficient approach to active learning by coupling uncertainty with principled blockwise distribution matching, providing a practical bridge to advanced optimizers without sacrificing accuracy.

REFERENCES

- Jordan T Ash, Chicheng Zhang, Akshay Krishnamurthy, John Langford, and Alekh Agarwal. Deep batch active learning by diverse, uncertain gradient lower bounds. *arXiv preprint arXiv:1906.03671*, 2019.
- Wonho Bae, Danica J Sutherland, and Gabriel L Oliveira. Uncertainty herding: One active learning method for all label budgets. In *ICLR*, 2025.
- Andrew An Bian, Baharan Mirzasoleiman, Joachim Buhmann, and Andreas Krause. Guaranteed non-convex optimization: Submodular maximization over continuous domains. In *Artificial Intelligence and Statistics*, pp. 111–120. PMLR, 2017.

- Zalán Borsos, Mojmir Mutny, and Andreas Krause. Coresets via bilevel optimization for continual learning and streaming. *Advances in neural information processing systems*, 33:14879–14890, 2020.
- Maximilian Böther, Abraham Sebastian, Pranjal Awasthi, Ana Klimovic, and Srikumar Ramalingam. On distributed larger-than-memory subset selection with pairwise submodular functions. *arXiv preprint arXiv:2402.16442*, 2024.
- Li Deng. The mnist database of handwritten digit images for machine learning research [best of the web]. *IEEE signal processing magazine*, 29(6):141–142, 2012.
- Edward Farhi, Jeffrey Goldstone, and Sam Gutmann. A quantum approximate optimization algorithm. *arXiv preprint arXiv:1411.4028*, 2014.
- Yarin Gal, Riashat Islam, and Zoubin Ghahramani. Deep bayesian active learning with image data. In *Proceedings of the 34th International Conference on Machine Learning - Volume 70, ICML’17*, pp. 1183–1192. JMLR.org, 2017.
- Fred Glover, Gary Kochenberger, and Yu Du. A tutorial on formulating and using qubo models. *arXiv preprint arXiv:1811.11538*, 2018.
- 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*, pp. 770–778, 2016.
- Tadashi Kadowaki and Hidetoshi Nishimori. Quantum annealing in the transverse ising model. *Physical Review E*, 58(5):5355, 1998.
- Gary Kochenberger, Jin-Kao Hao, Fred Glover, Mark Lewis, Zhipeng Lü, Haibo Wang, and Yang Wang. The unconstrained binary quadratic programming problem: a survey. *Journal of combinatorial optimization*, 28(1):58–81, 2014.
- Alex Krizhevsky, Geoffrey Hinton, et al. Learning multiple layers of features from tiny images. 2009.
- David D. Lewis and William A. Gale. A sequential algorithm for training text classifiers. In Bruce W. Croft and C. J. van Rijsbergen (eds.), *SIGIR ’94*, pp. 3–12, London, 1994. Springer London. ISBN 978-1-4471-2099-5.
- Dongyuan Li, Zhen Wang, Yankai Chen, Renhe Jiang, Weiping Ding, and Manabu Okumura. A survey on deep active learning: Recent advances and new frontiers. *IEEE Transactions on Neural Networks and Learning Systems*, 36(4):5879–5899, 2024.
- Stuart Lloyd. Least squares quantization in pcm. *IEEE transactions on information theory*, 28(2): 129–137, 1982.
- Andrew Lucas. Ising formulations of many np problems. *Frontiers in physics*, 2:5, 2014.
- Baharan Mirzasoleiman, Jeff Bilmes, and Jure Leskovec. Coresets for data-efficient training of machine learning models. In *International Conference on Machine Learning*, pp. 6950–6960. PMLR, 2020.
- Jishnu Mukhoti, Andreas Kirsch, Joost Van Amersfoort, Philip HS Torr, and Yarin Gal. Deep deterministic uncertainty: A new simple baseline. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pp. 24384–24394, 2023.
- Yuval Netzer, Tao Wang, Adam Coates, Alessandro Bissacco, Baolin Wu, Andrew Y Ng, et al. Reading digits in natural images with unsupervised feature learning. In *NIPS workshop on deep learning and unsupervised feature learning*, volume 2011, pp. 7. Granada, 2011.
- Robert Pinsler, Jonathan Gordon, Eric Nalisnick, and José Miguel Hernández-Lobato. Bayesian batch active learning as sparse subset approximation. *Advances in neural information processing systems*, 32, 2019.

- Ozan Sener and Silvio Savarese. Active learning for convolutional neural networks: A core-set approach. In *International Conference on Learning Representations*, 2018. URL <https://openreview.net/forum?id=H1aIuk-RW>.
- Burr Settles. Active learning literature survey. 2009.
- Claude E Shannon. A mathematical theory of communication. *ACM SIGMOBILE mobile computing and communications review*, 5(1):3–55, 2001.
- Atsushi Shimizu, Xiaouu Cheng, Christopher Musco, and Jonathan Weare. Improved active learning via dependent leverage score sampling. *arXiv preprint arXiv:2310.04966*, 2023.
- Samarth Sinha, Sayna Ebrahimi, and Trevor Darrell. Variational adversarial active learning. In *Proceedings of the IEEE/CVF international conference on computer vision*, pp. 5972–5981, 2019.
- Peiqi Wang, Yikang Shen, Zhen Guo, Matthew Stallone, Yoon Kim, Polina Golland, and Rameswar Panda. Diversity measurement and subset selection for instruction tuning datasets. *arXiv preprint arXiv:2402.02318*, 2024a.
- Yifeng Wang, Xueying Zhan, and Siyu Huang. Autoal: Automated active learning with differentiable query strategy search. *arXiv preprint arXiv:2410.13853*, 2024b.
- Han Xiao, Kashif Rasul, and Roland Vollgraf. Fashion-mnist: a novel image dataset for benchmarking machine learning algorithms. *arXiv preprint arXiv:1708.07747*, 2017.
- Donggeun Yoo and In So Kweon. Learning loss for active learning. In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, pp. 93–102, 2019.
- Xueying Zhan, Qingzhong Wang, Kuan-hao Huang, Haoyi Xiong, Dejing Dou, and Antoni B Chan. A comparative survey of deep active learning. *arXiv preprint arXiv:2203.13450*, 2022.
- Haizhong Zheng, Rui Liu, Fan Lai, and Atul Prakash. Coverage-centric coreset selection for high pruning rates. *arXiv preprint arXiv:2210.15809*, 2022.

LARGE MODEL USAGE STATEMENT

We used large language models for text translation and polishing.

A PROOF OF THEOREM 1

In this section we make precise the notion of a p -partial r -cover and present a complete proof of Theorem 1. We first recall the formal definition of p -partial r -cover. We then prove two-sided bounds that relate $\rho(r)$ to the empirical MMD cross term under a Gaussian kernel. This establishes a rigorous connection between RKHS-based distributional discrepancy and geometric coverage, clarifying why minimizing MMD promotes stronger coverage.

Definition 1 (p -partial r -cover). *Following (Zheng et al., 2022), given a metric space (X, d) and a probability measure μ on X (with density P_μ), a set $S \subset X$ is a p -partial r -cover if*

$$\int_X \mathbf{1}_{\bigcup_{x \in S} B_d(x, r)}(x) d\mu(x) = p,$$

where $B_d(x, r) = \{x' \in X : d(x, x') \leq r\}$ is the radius- r ball centered at x , and $\mathbf{1}_A$ denotes the indicator of a set A .

Proof. Let

$$A := \{x \in X : \text{dist}(x, S) \leq r\}, \quad |A| = |X| \rho(r). \quad (14)$$

Write the (standard) MMD cross average as

$$\delta_{\text{MMD}} = \frac{1}{|S| |X|} \sum_{s \in S} \sum_{x \in X} k(s, x) = \frac{1}{|X|} \sum_{x \in X} \left(\frac{1}{|S|} \sum_{s \in S} k(s, x) \right). \quad (15)$$

For any $x \in A$ there exists $s^* \in S$ with $\|x - s^*\| \leq r$, hence

$$\frac{1}{|S|} \sum_{s \in S} k(s, x) \geq \frac{1}{|S|} k(s^*, x) \geq \frac{1}{|S|} e^{-r^2/(2\sigma^2)}. \quad (16)$$

Averaging Eq.(16) over $x \in A$ and discarding the (nonnegative) contribution from $X \setminus A$ gives

$$\delta_{\text{MMD}} \geq \frac{|A|}{|X|} \cdot \frac{1}{|S|} e^{-r^2/(2\sigma^2)} = \frac{\rho(r)}{|S|} e^{-r^2/(2\sigma^2)}. \quad (17)$$

Therefore

$$\rho(r) \leq \delta_{\text{MMD}} |S| e^{r^2/(2\sigma^2)}. \quad (18)$$

For any $x \notin A$ we have $\text{dist}(x, S) > r$, so $k(s, x) \leq e^{-r^2/(2\sigma^2)}$ for all $s \in S$, while $k(s, x) \leq 1$ always. Hence

$$\frac{1}{|S|} \sum_{s \in S} k(s, x) \leq \begin{cases} 1, & x \in A, \\ e^{-r^2/(2\sigma^2)}, & x \notin A. \end{cases} \quad (19)$$

Averaging Eq.(19) over $x \in X$ yields

$$\delta_{\text{MMD}} \leq \frac{1}{|X|} \left(|A| \cdot 1 + (|X| - |A|) e^{-r^2/(2\sigma^2)} \right) = \rho(r) + (1 - \rho(r)) e^{-r^2/(2\sigma^2)}. \quad (20)$$

Solving Eq.(20) for $\rho(r)$ gives

$$\rho(r) \geq \frac{\delta_{\text{MMD}} - e^{-r^2/(2\sigma^2)}}{1 - e^{-r^2/(2\sigma^2)}}. \quad (21)$$

Combining Eq.(18) and Eq.(21) proves Eq.(7). $\square$

B FROM MMD MINIMIZATION TO QUBO

We briefly show how the cardinality-constrained MMD objective becomes a single-matrix QUBO. From Eq.(22), (i) add the quadratic penalty $\lambda(\mathbf{1}^\top m - k)^2$ to enforce $\mathbf{1}^\top m = k$, and (ii) use the binary identity $m_i^2 = m_i$ to absorb the linear term into the diagonal, yielding a standard QUBO with matrix $\hat{Q}$. For sufficiently large λ , minimizers coincide. We also give element-wise coefficients of $\hat{Q}$ for implementation.

Let $K \in \mathbb{R}^{n \times n}$ be the Gram matrix with $K_{ij} = k(x_i, x_j)$ and $m \in \{0, 1\}^n$ encode the selected subset with $\mathbf{1}^\top m = k$. The empirical MMD objective between the subset and the full pool reduces to

$$\min_{m \in \{0,1\}^n} \underbrace{\frac{1}{k^2} m^\top K m}_{\text{quadratic}} - \underbrace{\frac{2}{nk} \mathbf{1}^\top K m}_{\text{linear}} \quad \text{s.t.} \quad \mathbf{1}^\top m = k. \quad (22)$$

Step 1 (penalize the cardinality constraint). Add a quadratic penalty with $\lambda > 0$:

$$\lambda(\mathbf{1}^\top m - k)^2 = \lambda m^\top J m - 2k\lambda \mathbf{1}^\top m + \lambda k^2, \quad J := \mathbf{1}\mathbf{1}^\top.$$

Dropping the constant λk^2 , we obtain the unconstrained binary quadratic objective

$$\min_{m \in \{0,1\}^n} m^\top Q m + q^\top m + \text{const}, \quad Q := \frac{1}{k^2} K + \lambda J, \quad q := -\frac{2}{nk} K \mathbf{1} - 2k\lambda \mathbf{1}. \quad (23)$$

For sufficiently large λ , any minimizer of Eq.(23) satisfies $\mathbf{1}^\top m = k$ and thus solves Eq.(22).

Step 2 (absorb the linear term into the quadratic). For binary m , $m_i^2 = m_i$, hence $q^\top m = \sum_i q_i m_i = \sum_i q_i m_i^2 = m^\top \text{Diag}(q) m$. Define

$$\hat{Q} := Q + \text{Diag}(q) = \frac{1}{k^2} K + \lambda J + \text{Diag}\left(-\frac{2}{nk} K \mathbf{1} - 2k\lambda \mathbf{1}\right).$$

Then we obtain a standard single-matrix QUBO:

$$\min_{m \in \{0,1\}^n} m^\top \hat{Q} m \quad (24)$$

and Eq.(24) and Eq.(22) have the same minimizers (constants omitted). This QUBO can be handled by classical QUBO solvers or quantum Ising backends; the detailed equivalence proof and choices of λ are in Appendix B.

Element-wise coefficients of $\hat{Q}$. Let $u := K \mathbf{1}$. Then

$$\hat{Q}_{ij} = \begin{cases} \frac{1}{k^2} K_{ij} + \lambda, & i \neq j, \\ \frac{1}{k^2} K_{ii} + \lambda - \frac{2}{nk} u_i - 2k\lambda, & i = j. \end{cases}$$

Here J contributes λ to *all* entries, including the diagonal; the linear term contributes only to the diagonal via $\text{Diag}(q)$.

C PROOF OF THEOREM2

We bound the global RKHS discrepancy between the selected set and the pool by separating it into (a) within-block mismatches and (b) a quota-mismatch term that measures how far the selection weights α deviate from the block proportions w . Using a two-term inequality with a tunable constant, convexity, and bounded-kernel arguments, we obtain the blockwise approximation bound in Eq.(31); when quotas are proportional ($\alpha = w$), it simplifies to Eq.(32). The details follow.

We first note the blockwise decompositions

$$\mu_S = \sum_{b=1}^B \alpha_b \mu_{S_b}, \quad \mu_X = \sum_{b=1}^B w_b \mu_{X_b}, \quad (25)$$

which imply

$$\mu_S - \mu_X = \underbrace{\sum_{b=1}^B \alpha_b (\mu_{S_b} - \mu_{X_b})}_{\triangleq U} + \underbrace{\sum_{b=1}^B (\alpha_b - w_b) \mu_{X_b}}_{\triangleq V}. \quad (26)$$

(i) Two-term split with a tunable constant. For any $\eta > 0$, the parallelogram-type inequality

$$\|U + V\|^2 \leq (1 + \eta)\|U\|^2 + \left(1 + \frac{1}{\eta}\right)\|V\|^2$$

holds. Taking $\eta = 1$ recovers the bound used in the main text:

$$\|\mu_S - \mu_X\|^2 \leq 2\|U\|^2 + 2\|V\|^2. \quad (27)$$

This shows where the factor 2 comes from and also allows a trade-off via η when one term is known to be small.

(ii) First term via Jensen/convexity. Because $\sum_b \alpha_b = 1$ and $\|\cdot\|^2$ is convex,

$$\left\| \sum_{b=1}^B \alpha_b (\mu_{S_b} - \mu_{X_b}) \right\|^2 \leq \sum_{b=1}^B \alpha_b \|\mu_{S_b} - \mu_{X_b}\|^2 = \sum_{b=1}^B \alpha_b \text{MMD}^2(S_b, X_b). \quad (28)$$

(iii) Second term via bounded kernels. Assume $\sup_x k(x, x) \leq \kappa^2$. Then $\|\phi(x)\| \leq \kappa$ and hence

$$\|\mu_{X_b}\| = \left\| \frac{1}{|X_b|} \sum_{x \in X_b} \phi(x) \right\| \leq \frac{1}{|X_b|} \sum_{x \in X_b} \|\phi(x)\| \leq \kappa. \quad (29)$$

By the triangle inequality,

$$\left\| \sum_{b=1}^B (\alpha_b - w_b) \mu_{X_b} \right\| \leq \sum_{b=1}^B |\alpha_b - w_b| \|\mu_{X_b}\| \leq \kappa \|\alpha - w\|_1. \quad (30)$$

Combining Eq.(27)– Eq.(30) yields the general bound

$$\text{MMD}^2(S, X) \leq 2 \sum_{b=1}^B \alpha_b \text{MMD}^2(S_b, X_b) + 2\kappa^2 \|\alpha - w\|_1^2. \quad (31)$$

(iv) Proportional quotas as a corollary. When $\alpha = w$ (quotas proportional to block sizes), the mismatch term vanishes and Eq.(31) simplifies to

$$\text{MMD}^2(S, X) \leq 2 \sum_{b=1}^B w_b \text{MMD}^2(S_b, X_b) \leq 2 \max_b \text{MMD}^2(S_b, X_b), \quad (32)$$

which is exactly the theorem stated in the merged main text.

(v) Optional variants. If additional structure is available, one may tighten the mismatch term:

- Using $\|\sum_b c_b \mu_{X_b}\| \leq (\sum_b c_b^2)^{1/2} (\sum_b \|\mu_{X_b}\|^2)^{1/2} \leq \kappa \sqrt{B} \|\alpha - w\|_2$ gives an alternative bound $2\kappa^2 B \|\alpha - w\|_2^2$, which can be sharper or looser than the ℓ_1 bound depending on B and the sparsity of $\alpha - w$.
- If blocks are very compact so that $\|\mu_{X_b} - \mu_X\| \leq \Delta$ for all b , then writing $V = \sum_b (\alpha_b - w_b)(\mu_{X_b} - \mu_X)$ yields $\|V\| \leq \Delta \|\alpha - w\|_1$, replacing κ by the typically smaller dispersion Δ .

(vi) Equality and interpretation. If $\alpha = w$ and each block attains $\mu_{S_b} = \mu_{X_b}$ (perfect local matching), then the right-hand side of Eq.(32) is 0, hence $\text{MMD}^2(S, X) = 0$. Thus compact clustering plus proportional quotas makes the global error a convex combination of local errors, explaining why blockwise MMD optimization closely tracks the full-pool objective.

D ISING MODEL

This subsection has two parts. First, we recall the standard, closed-form mapping from a binary QUBO to an Ising model via the change of variables $x_i = (s_i + 1)/2$, yielding explicit expressions for the pairwise couplers J_{ij} , local fields h_i , and the constant shift. Second, we address a practical issue on current quantum/annealing hardware: coefficients are encoded with limited bitwidth and dynamic range. In many AL objectives the cardinality penalty aggregates into large local fields h_i whose scale can dwarf the pairwise couplers J_{ij} ; after global rescaling/quantization this compresses the effective resolution of J_{ij} . We therefore describe a *multi-auxiliary* scheme that redistributes overly large data-spin fields into several couplers to a small set of ancilla spins, so that each coupling sits on the same scale as typical J_{ij} . This improves numerical conditioning and retains more useful precision for J in QUBO→Ising deployments.

D.1 STANDARD QUBO TO ISING MAPPING

Consider a QUBO in the upper-triangular form

$$\min_{x \in \{0,1\}^n} E_{\text{QUBO}}(x) = \sum_{i \leq j} Q_{ij} x_i x_j + C_0, \quad (33)$$

where linear terms have been absorbed into the diagonal (Q_{ii}). Introduce Ising spins $s \in \{-1, 1\}^n$ via $x_i = (s_i + 1)/2$. Using

$$x_i x_j = \frac{1}{4} (1 + s_i + s_j + s_i s_j) \quad (i < j), \quad x_i = \frac{1}{2} (1 + s_i),$$

we obtain an Ising Hamiltonian (up to an additive constant)

$$\min_{s \in \{-1,1\}^n} E_{\text{Ising}}(s) = - \sum_{i < j} J_{ij} s_i s_j - \sum_i h_i s_i + C, \quad (34)$$

with coefficients

$$J_{ij} = -\frac{Q_{ij}}{4} \quad (i < j), \quad h_i = -\frac{Q_{ii}}{2} - \frac{1}{4} \sum_{j \neq i} Q_{ij}, \quad (35)$$

and constant $C = C_0 + \frac{1}{4} \sum_{i < j} Q_{ij} + \frac{1}{2} \sum_i Q_{ii}$ (which does not affect minimizers). Equations Eq.(33) and Eq.(34) therefore have the same minimizers after the change of variables $x_i = (s_i + 1)/2$.

Multi-auxiliary (precision-aware) construction We start from the Ising form

$$E(s) = - \sum_{i < j} J_{ij} s_i s_j - \sum_{i=1}^n h_i s_i, \quad s_i \in \{-1, 1\}. \quad (36)$$

When the bitwidth is limited, very large $|h_i|$ (e.g., from a cardinality penalty) dominate the global scale used for coefficient normalization/quantization, thereby squeezing the dynamic range available to represent J_{ij} . To alleviate this, we split each large field into several couplers to auxiliary spins so that every individual coupling magnitude is comparable to a chosen J -scale.

Concretely, introduce G ancilla spins $a_g \in \{-1, 1\}$ and nonnegative weights w_{ig} such that

$$\sum_{g=1}^{G_i} w_{ig} = |h_i|, \quad w_{ig} \approx \Delta_J, \quad (37)$$

where $G_i = \max(1, \lceil |h_i|/(\Delta_J) \rceil)$ and Δ_J is a reference coupler scale (e.g., $\Delta_J = \text{median}_{i < j} |J_{ij}|$ or $\max_{i < j} |J_{ij}|$). Let $\sigma_i = \text{sign}(h_i)$. We implement the following bounded-scale Ising model on hardware:

$$E_{\text{split}}(s, a) = - \sum_{i < j} J_{ij} s_i s_j - \sum_{g=1}^G \left(\sum_{i=1}^n \sigma_i w_{ig} s_i \right) a_g - \sum_{g=1}^G B_g a_g, \quad (38)$$

with modest ancilla biases B_g chosen by a simple rule such as

$$B_g = (1 + \epsilon) \sum_{i=1}^n w_{ig} \quad \text{with} \quad 0 < \epsilon \ll 1, \quad (39)$$

and then jointly rescaled with $\{J_{ij}\}$ to the device’s representable range. By construction, each coupling $|\sigma_i w_{ig}|$ is of the same order as Δ_J , so after the single global rescaling/quantization step the J_{ij} entries retain substantially more effective resolution. In practice this “field-to-coupler splitting” acts as a precision-aware reparameterization: it reduces dynamic-range imbalance between h and J without altering the data–data couplers $\{J_{ij}\}$.

E IMPLEMENTATION DETAILS

We benchmark on five image–classification datasets, deliberately using relatively small batch sizes b to emphasize the small–batch regime where diversity matters. Table 2 summarizes the datasets and our default AL settings. We largely follow prior configurations, but intentionally adopt relatively small batch sizes b to highlight the importance of diversity in the small–batch regime. Unless otherwise noted, the backbone is ResNet-18 (He et al., 2016). Our DAL framework is built on the evaluation setup of (Zhan et al., 2022); please refer to their paper for additional architectural and training details.

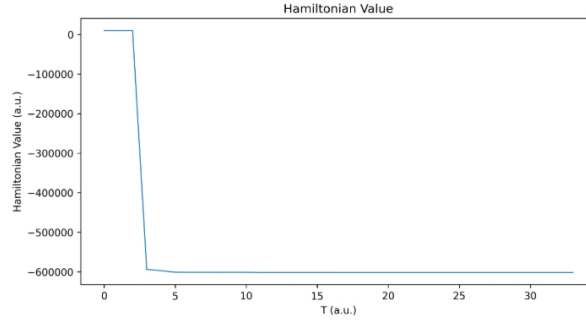
i is the size of initial labeled pool, u is the size of unlabeled data pool, t is the size of testing set, b is the per-round batch size, Q is the total query budget, C is the number of categories, and e is the number of epochs used to train the basic classifier in each AL round. Each algorithm is run three times per dataset to report the mean and variance. $M_{\max}$ is the maximum block size and τ is the candidate ratio (fraction of unlabeled pool sampled as candidates). For fairness, we set the number of training epochs for *Uherding* to match the other methods, while keeping its original learning rate 0.025; empirically, reducing it to 0.001 (the rate used by other methods) led to much slower accuracy improvements. Because the public *Uherding* codebase does not provide configurations for *SVHN*, *MNIST*, and *FashionMNIST*, we did not reproduce *Uherding* on these three datasets. The head-to-head comparison between the quantum solver and the greedy algorithm was conducted with batch size $b = 100$ and candidate ratio $\tau = 0.06$, reflecting the nontrivial overhead of invoking current quantum hardware and chosen to conserve computational resources.

When blocks are used, we fix the maximum block size $M_{\max}$, adopt proportional quotas ($\alpha_b = w_b$), and use the same candidate ratio τ to form UCPs. The kernel function we adopt is the Gaussian kernel function. Kernel bandwidths are set per block via the median heuristic; the solver for the MMD subproblem is identical across runs. Each method is run three times per dataset with different seeds.

Dataset	i	u	t	b	Q	C	e	$M_{\max}$	τ
<i>MNIST</i>	100	59,500	10,000	50	1,000	10	20	5,000	0.1
<i>FashionMNIST</i>	500	59,500	10,000	50	5,000	10	20	5,000	0.1
<i>SVHN</i>	500	72,757	26,032	50	5,000	10	20	5,000	0.1
<i>CIFAR10</i>	1,000	49,000	10,000	100	6,000	10	30	5,000	0.1
<i>CIFAR100</i>	1,000	49,000	10,000	100	10,000	100	40	5,000	0.1

Table 2: Datasets used in comparative experiments.

Figure 4 shows the measured expectation of the problem Hamiltonian during a single annealing run. The characteristic steep initial descent reflects the transition from the initial driver Hamiltonian to the problem Hamiltonian, during which the system rapidly explores the energy landscape and converges toward low-energy regions. The long, flat plateau that follows indicates the system has settled into a metastable state within a low-energy basin, which corresponds to the solution subset returned by the annealer. In our encoding, a lower Hamiltonian expectation corresponds directly to a smaller MMD objective. Therefore, the observed trajectory confirms that the quantum annealer successfully progresses to a high-quality solution.

Figure 4: **Hamiltonian evolution.**

F ADDITIONAL PAIRWISE PLOTS (ABLATIONS)

In Figure 5 we present *pairwise* comparisons that isolate the effect of our diversity module: for each selector (Entropy, BALD, LPL) we plot its uncertainty-only variant against its “EBDAL” counterpart under identical backbones, training schedules, and budgets. The only change within each pair is the addition of blockwise candidate filtering and MMD-based selection.

Across datasets, the EBDAL curves consistently *lift and smooth* the accuracy–budget trajectories relative to their baselines, with the largest gains on harder settings such as CIFAR-100. Improvements are already visible at small budgets—where uncertainty-only methods tend to sample clustered, redundant points—and remain evident as labeling proceeds on CIFAR-10, SVHN, MNIST, and FashionMNIST.

Overall, these plots confirm that coupling uncertainty with blockwise diversity reduces small-batch redundancy and improves label efficiency in a method-agnostic way.

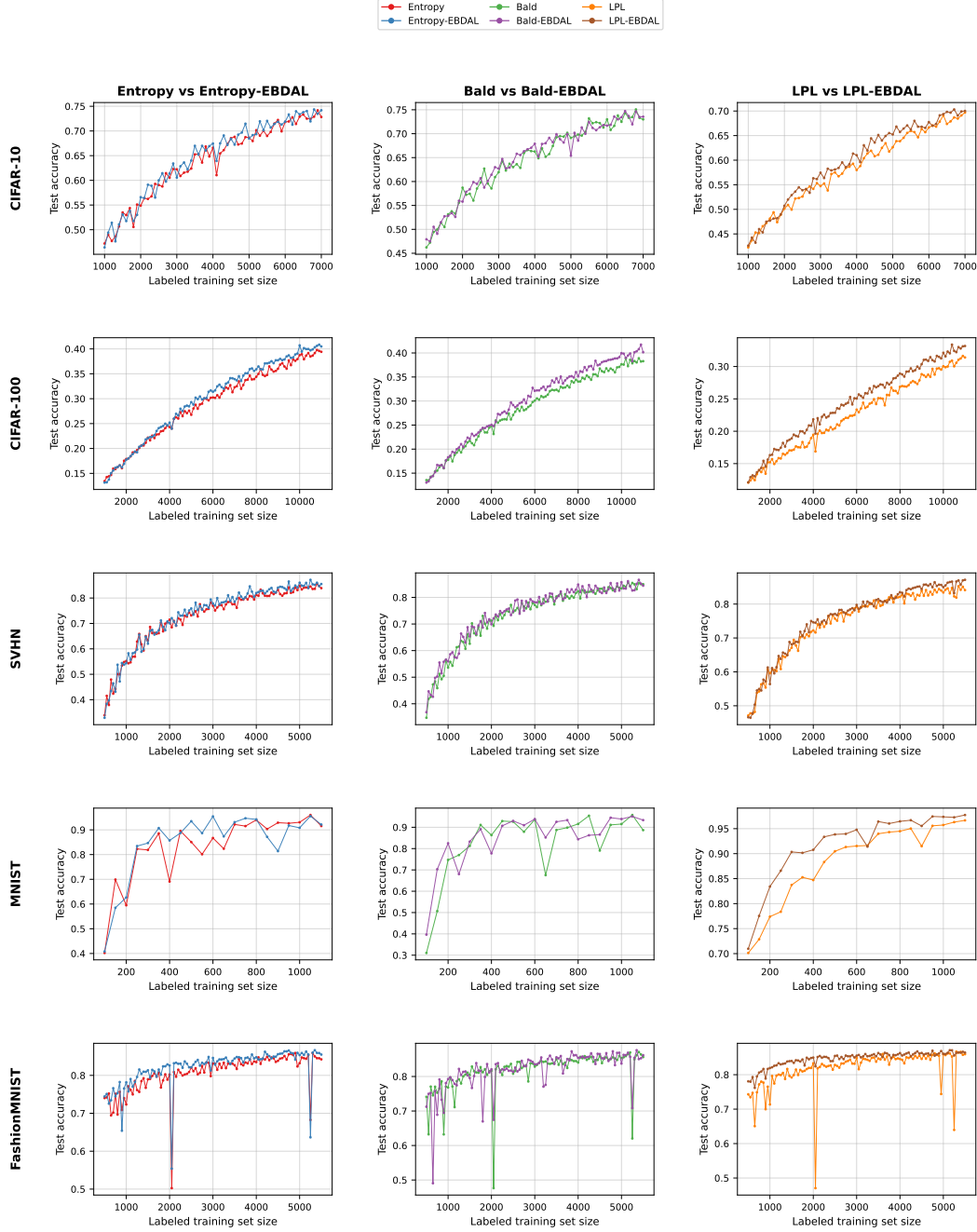


Figure 5: **Pairwise comparisons of uncertainty-only vs. EBDAL variants.** Columns: (left) Entropy vs. Entropy-EBDAL; (middle) BALD vs. BALD-EBDAL; (right) LPL vs. LPL-EBDAL. Rows: CIFAR-10, CIFAR-100, SVHN, MNIST, FashionMNIST. EBDAL smooths and elevates the curves at the same budget.