

Cleaning the Pool: Progressive Filtering of Unlabeled Pools in Deep Active Learning

Denis Huseljic Marek Herde Lukas Rauch Paul Hahn Bernhard Sick
Intelligent Embedded Systems, University of Kassel
{firstname.lastname}@uni-kassel.de

Abstract

Existing active learning (AL) strategies capture fundamentally different notions of data value, e.g., uncertainty or representativeness. Consequently, the effectiveness of strategies can vary substantially across datasets, models, and even AL cycles. Committing to a single strategy risks suboptimal performance, as no single strategy dominates throughout the entire AL process. We introduce REFINE, an ensemble AL method that combines multiple strategies without knowing in advance which will perform best. In each AL cycle, REFINE operates in two stages: (1) Progressive filtering iteratively refines the unlabeled pool by considering an ensemble of AL strategies, retaining promising candidates capturing different notions of value. (2) Coverage-based selection then chooses a final batch from this refined pool, ensuring all previously identified notions of value are accounted for. Extensive experiments across 6 classification datasets and 3 foundation models show that REFINE consistently outperforms individual strategies and existing ensemble methods. Notably, progressive filtering serves as a powerful preprocessing step that improves the performance of any individual AL strategy applied to the refined pool, which we demonstrate on an audio spectrogram classification use case. Finally, the ensemble of REFINE can be easily extended with upcoming state-of-the-art AL strategies.

1. Introduction

Pretrained foundation models have become the backbone of modern machine learning, offering powerful general-purpose representations [30]. Yet, adapting these models to downstream tasks still often requires large amounts of annotated data. Obtaining annotations is costly and can quickly become a bottleneck in many applications. Active learning (AL) aims to mitigate this through an intelligent selection of instances for annotation [29]. Considering a pool-based batch AL setting, AL strategies query multiple instances per

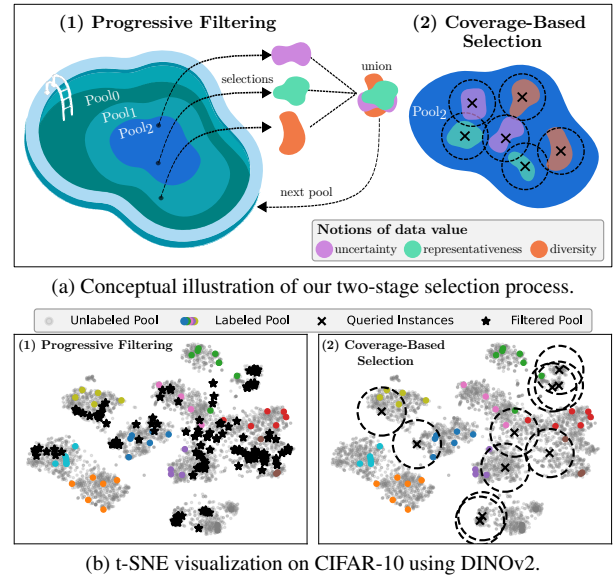


Figure 1. REFINE’s two-stage selection process. Stage 1: Progressive filtering refines the unlabeled pool through multiple rounds, where each round outputs the union of instances selected by each strategy as the next round’s input pool. Stage 2: Coverage-based selection then chooses the final batch from this refined pool, ensuring a selection that accounts for diverse notions of value.

cycle from a large unlabeled pool, gradually expanding the training set while reducing annotation costs.

One central challenge in AL is identifying which instances maximize model performance. In recent years, much progress has been made in developing various selection strategies [1, 2, 4]. However, recent benchmarks show that their effectiveness varies substantially across models, domains, and even cycles of the AL process [19, 21, 26]. Moreover, choosing a suboptimal strategy can even yield worse performance than simple random sampling [21]. This variability poses a significant challenge for practitioners seeking to select an appropriate strategy for their downstream task, especially since AL is essentially a one-shot problem with little opportunity for trial and error.

Part of the challenge is that different selection strategies capture fundamentally different notions of data value, *e.g.*, uncertainty-based strategies target difficult instances while representativeness-based strategies seek high-density ones. Each perspective offers reasonable insights, yet no single strategy dominates across the entire AL process, as instances’ impact on model performance strongly varies per cycle [4, 11]. As the model learns, the optimal selection criterion shifts: a strategy that excels in one AL cycle may become suboptimal in the next. This means that committing to any single strategy requires accepting its blind spots. While finding a universally optimal strategy is challenging, we presume that combining multiple complementary strategies in an *ensemble AL method* provides a more robust alternative. While leveraging collective strengths, we reduce the risk of committing to any potentially suboptimal heuristic.

Our approach. We propose REFINE, an ensemble AL method that avoids committing to a single strategy. REFINE builds on the core insight that combining multiple complementary strategies, each considering its own notion of data value, can yield an improved selection that captures diverse perspectives of what makes an instance valuable.

Specifically, in each AL cycle, REFINE operates in two stages (Fig. 1): **(1) Progressive filtering** iteratively refines the unlabeled pool through multiple rounds. In each round, we apply all strategies from the ensemble on the current unlabeled pool and take the union of the selected instances, retaining batches that at least one strategy considers valuable. Repeating this process yields an high-value candidate pool, *i.e.*, it reflects complementary notions of data value while filtering out uninformative instances that provide no value according to any strategy. This pool cleaning process is a powerful preprocessing step on its own, already considerably improving the performance of individual strategies. **(2) Coverage-based selection** then selects the final batch from the refined pool by choosing the instances that, when combined with the already labeled pool, maximize coverage over the refined pool. Since progressive filtering has already identified high-value candidates, this selection ensures the batch is accounting for the various notions of data value captured in the filtering stage.

Our contributions can be summarized as follows:

- We introduce *progressive filtering*, an approach to iteratively refine the unlabeled pool by retaining valuable instances while cleaning out uninformative ones, thereby reducing pool size and improving pool quality for downstream selection.
- We propose REFINE, an *ensemble AL method* that combines progressive filtering with a final coverage-based selection to leverage complementary notions of data value captured by multiple AL strategies in the refined pool.
- We provide a *theoretical analysis* of progressive filtering, establishing bounds on its ability to retain high-value in-

stances and remove uninformative ones.

- We demonstrate that REFINE *outperforms* individual strategies and existing ensemble AL methods across 6 image classification datasets and 3 foundation models. Additionally, we validate the practicality of progressive filtering on an audio spectrogram classification use case.

2. Related Work

AL strategies can be categorized into at least one of two criteria: uncertainty or representativeness. *Margin* [29] focuses on uncertainty by selecting the top- k instances with the smallest difference between the two most probable classes. *TypiClust* [12] prioritizes representativeness via high-density instances, employing k -MEANS for batch diversity. Similarly, *BADGE* [2] targets uncertainty through large gradient updates, utilizing k -MEANS++ to diversify the batch. *BAIT* [1] greedily selects instances maximizing Fisher information. Both *AlfaMix* [23] and *DropQuery* [10] use perturbations (mixed features or dropout) to filter instances based on prediction variability, followed by k -MEANS clustering. Finally, *MaxHerding* [3] greedily maximizes generalized coverage via feature-based kernel similarities, which *UHerding* [4] extends by incorporating an uncertainty factor.

Ensemble AL methods use multiple AL strategies by either *combining them* or *choosing when to use which*. Although *ALBL* [14], *COMB* [5], and *DUAL* [8] fit this regime, we do not consider them as they are single instance strategies whose adaptation to the batch setting is not straightforward. *SelectAL* [11] chooses between low- and high-budget strategies by evaluating which one identifies the most valuable batch of the current labeled pool. However, consistent with Bae et al. [4], we found it yields inconsistent performance and is difficult to reproduce. *TAILOR* [31] frames the choice of AL strategies as a multi-armed bandit problem, learning a probability distribution used to sample which strategies propose instances. As TAILOR uses a class-balance-based reward to update this distribution, it can handle class-imbalanced settings effectively. *TCM* [9] initially applies TypiClust for representativeness before transitioning to Margin for uncertainty, determining the switch point via a budget-dependent heuristic. This design is based on empirical findings that representative and diverse instances perform well on low budgets, while uncertainty is well-suited for high budgets [11].

Unlike previous ensemble AL methods that attempt to learn or plan to switch between strategies, we propose *progressive filtering*, a robust and training-free method. Instead of selecting a single strategy, this mechanism iteratively refines the unlabeled pool to retain high-value candidates from all available AL strategies. Furthermore, unlike strategies such as DropQuery or AlfaMix that filter in one step based on a single fixed heuristic, we generalize this concept

by filtering across multiple rounds using diverse notions of data value. We provide both a theoretical analysis and empirical validation for the effectiveness of this concept.

3. Notation

We consider a pool-based batch AL setting for classification. Let $x \in \mathcal{X}$ be an instance and $y \in \mathcal{Y} = \{1, \dots, K\}$ denote its label, where K is the number of classes. Further, let $\mathcal{U}_t \subset \mathcal{X}$ be a large unlabeled pool and $\mathcal{L}_t \subset \mathcal{X} \times \mathcal{Y}$ be the labeled pool at cycle t . Let $\mathcal{S} = \{s_1, \dots, s_M\}$ denote the set of M selection strategies in the ensemble. At $t = 0$, we initialize $\mathcal{L}_0$ by randomly sampling b instances. Then, we perform a total of A AL cycles, selecting b instances to label in each cycle. The total labeling budget is $B = b + A \cdot b$. We consider foundation models (or backbones) consisting of feature extractor $h^\phi : \mathcal{X} \rightarrow \mathbb{R}^D$ and classification head $g^\theta : \mathbb{R}^D \rightarrow \mathbb{R}^K$, where ϕ and θ are fixed and trainable parameters, respectively. Hence, our model is a function $f = (g^\theta \circ h^\phi)(x)$ mapping an instance to the logit space.

4. Method

In this section, we detail REFINE with its two-stage selection process. We first introduce progressive filtering to refine the unlabeled pool for an improved downstream selection. We then describe how to effectively sample from this refined pool, accounting for the added value of each strategy. Afterward, we analyze theoretical properties of progressive filtering, demonstrating that it preserves high-value candidates while filtering out uninformative ones.

4.1. Progressive Filtering

We define the *value* of an instance as its potential to improve model performance. What constitutes value depends on a strategy’s perspective, such as reducing uncertainty or ensuring representativeness. An instance is considered *collectively valuable* if it is ranked highly by at least one strategy, representing value from at least one perspective.

In each cycle, progressive filtering iteratively cleans the unlabeled pool to identify a subset of candidates that are collectively valuable across multiple selection strategies. Our core idea is to let selection strategies act as experts, each providing a distinct perspective on data value. When an expert considers a batch to be valuable, we include it in the refined pool. Through multiple iterations, this process progressively prioritizes high-value instances while filtering out uninformative ones. Intuitively, while experts may not agree on what is valuable, instances never selected by any expert are likely to be uninformative.

Formally, given the unlabeled pool $\mathcal{U}_t$ at cycle t and the set of M selection strategies $\mathcal{S}$, our goal is to obtain a candidate pool $\mathcal{C}_R \subset \mathcal{U}_t$ through R rounds. In each round $r \in \{1, \dots, R\}$, each of the M strategies select J batches

Algorithm 1 Progressive Filtering

Require: Unlabeled pool $\mathcal{U}_t$, selection strategies $\mathcal{S} = \{s_1, \dots, s_M\}$, number of rounds R , number of batches per strategy J , batch size b , sampling ratio α

```

1:  $\mathcal{C}_0 \leftarrow \mathcal{U}_t$ 
2: for  $r \in \{1, \dots, R\}$  do
3:    $\mathcal{C}_r \leftarrow \emptyset$ 
4:   for  $m \in \{1, \dots, M\}$  do
5:     for  $j \in \{1, \dots, J\}$  do
6:        $\mathcal{U}_{\text{sample}} \leftarrow \text{SubSample}(\mathcal{C}_{r-1}, \alpha \cdot |\mathcal{C}_{r-1}|)$ 
7:        $\mathcal{C}_r \leftarrow \mathcal{C}_r \cup s_m(\mathcal{U}_{\text{sample}}, b)$ 
8:     end for
9:   end for
10: end for
11: return  $\mathcal{C}_R$ 
```

from the current refined pool $\mathcal{C}_{r-1}$. By selecting multiple batches ($J > 1$), we prevent the pool $\mathcal{C}_r$ from shrinking too rapidly and give each strategy multiple chances to identify valuable instances. To ensure computational tractability and to enable deterministic strategies to produce non-identical batches, each strategy s_m selects from a randomly sampled subset of $\mathcal{C}_{r-1}$. The refined pool for round r is then the union of all selected batches:

$$\mathcal{C}_r = \bigcup_{m=1}^M \bigcup_{j=1}^J s_m(\text{SubSample}(\mathcal{C}_{r-1}, \alpha \cdot |\mathcal{C}_{r-1}|), b), \quad (1)$$

where $\mathcal{C}_0 = \mathcal{U}_t$, $\alpha \in (0, 1)$ is the sample ratio, and $\text{SubSample}(\cdot)$ draws instances uniformly without replacement. After R rounds, the candidate pool $\mathcal{C}_R$ serves as the refined pool from which the final batch for annotation is selected. The approach is summarized in Algorithm 1.

Design Rationale. Three key design choices underpin progressive filtering. *Union Over Intersection:* While intersection might seem appealing for finding consensus instances, it has two critical limitations: (i) it discards instances that are uniquely valuable to specific strategies, reducing diversity, and (ii) it can result in (near-)empty sets when strategies have little overlap. The union operation preserves all potentially valuable instances while filtering uninformative instances through repetition. *Iterative Refinement:* A single filtering round would simply concatenate all strategies’ selections, providing only a slight reduction in uninformative instances. The power of progressive filtering emerges through multiple rounds: instances must be repeatedly selected to remain in the pool. Since each round operates on the output of the previous round, uninformative instances are unlikely to be consistently ranked high. Survival of instances through rounds requires consistent recognition of value. *Sample Ratio:* The sample ratio α controls the stochasticity of the filtering process. Setting α close to 1 makes the selection process more deterministic, which can

cause the pool to shrink too rapidly. Reducing α , we introduce controlled randomness that: (i) enables deterministic strategies to produce diverse batches, (ii) amplifies diversity by forcing exploration across different regions of the candidate space, and (iii) creates robust filtering where valuable instances have multiple chances to be selected. Furthermore, a smaller value for α reduces the computational and memory cost of selection strategies by limiting the pool size. In our experiments, we found that almost any value of α works well, provided it is not too small (*cf.* Sec. 5).

Practical Considerations. We now discuss key implementation choices for applying progressive filtering. *Strategy Set Composition:* The composition of the set $\mathcal{S}$ represents a key design choice. Rather than attempting to find an optimal subset of strategies, which might be dataset- or model-dependent, we adopt a simple and principled approach: employing all individual batch AL strategies available in our benchmark in Sec. 5. Progressive filtering places no restrictions on the types of AL strategies in its ensemble, requiring only that each can produce a batch of candidate instances. By incorporating all available perspectives, our method naturally captures diverse notions of value without manual tuning and prior knowledge. Moreover, this also makes progressive filtering easily expandable, allowing the ensemble to be extended with novel AL strategies that may offer new perspectives on instances’ values. *Minimum Pool Size:* We enforce a minimum pool size equal to the batch size ($|\mathcal{C}_R| \geq b$) to ensure that the subsequent coverage-based selection has sufficient instances for final selection. Moreover, the parameter α also mitigates this risk by introducing randomness into selections, preventing excessive shrinkage of $\mathcal{C}_R$. *Hyperparameters:* We set $\alpha = 0.4$, $R = 5$, and $J = 10$ based on a grid search on a validation split of CIFAR-10 and fix these values across all experiments. We further analyze their impact in Sec. 5.

Computational and Memory Efficiency. Progressive filtering incurs additional computational overhead compared to standard batch selection, but this cost is justified by improved selection quality. The computational complexity per round is $\mathcal{O}(M \cdot J \cdot C_{\max}(\alpha \cdot |\mathcal{C}_{r-1}|))$, where $C_{\max}(\mathcal{C})$ is the cost of applying the most computationally expensive strategy on pool $\mathcal{C}$. Two factors help control costs: (i) the sampling ratio $\alpha < 1$ reduces the pool size each strategy must evaluate, and (ii) $\mathcal{C}_r$ is guaranteed to be a subset of $\mathcal{C}_{r-1}$, ensuring the pool size decreases monotonically. This design also addresses a common memory bottleneck. Many AL strategies do not scale to extensively large pools, often requiring a fixed subsample of $\mathcal{U}_t$ to be applicable. In contrast, by repeatedly using subsets of size $\alpha \cdot |\mathcal{C}_{r-1}|$, progressive filtering can explore the entire unlabeled pool without the large memory footprint of evaluating the entire unlabeled pool at once. Overall, while progressive filtering incurs overhead, it is a reasonable trade-off for the perfor-

mance gains, especially since the execution of AL strategies can be parallelized in practice. Assuming M concurrent processes, the cost of progressive filtering effectively reduces to $\mathcal{O}(J \cdot C_{\max}(\alpha \cdot |\mathcal{C}_{r-1}|))$ per round.

4.2. Coverage-Based Selection

Given the refined pool $\mathcal{C}_R$ from progressive filtering, we now address how to select the final batch for annotation. At this stage, $\mathcal{C}_R$ consists primarily of valuable instances identified through collective assessment by multiple strategies. The remaining challenge is to select a batch that effectively leverages these diverse perspectives.

We select instances that, when added to $\mathcal{L}_t$, best represent the distribution of $\mathcal{C}_R$. Specifically, we seek a batch $\mathcal{B}^* \subset \mathcal{C}_R$ that maximizes coverage [3, 4] over $\mathcal{C}_R$:

$$\mathcal{B}^* = \arg \max_{\mathcal{B} \subset \mathcal{C}_R, |\mathcal{B}|=b} \mathbb{E}_{\mathbf{x}} \left[\max_{\mathbf{x}' \in (\mathcal{L}_t \cup \mathcal{B})} k(\mathbf{x}, \mathbf{x}') \right], \quad (2)$$

where $k(\mathbf{x}, \mathbf{x}')$ is a kernel measuring the similarity between instances. Intuitively, maximizing Eq. (2) ensures that every instance in $\mathcal{C}_R$ is well-represented by at least one instance in $\mathcal{B} \cup \mathcal{L}_t$. This way, we capture diverse notions of value through progressive filtering in the final selection. Considering foundation models, Eq. (2) is maximized by applying the kernel to features $h^\phi(\mathbf{x})$ rather than to the raw instances.

There are several existing AL strategies that explicitly optimize for coverage, typically over the entire unlabeled pool [3, 4, 12, 28]. In our implementation, we adopt UHerding [4], which extends MaxHerding [3] by greedily maximizing coverage with an additional uncertainty factor that balances exploration and exploitation. We found this factor to be beneficial as most of the strategies in $\mathcal{S}$ already focus on representativeness. Consequently, the final pool $\mathcal{C}_R$ has a stronger focus on representative instances. However, any coverage-based selection strategy, such as TypiClust [12], could be employed.

Design Rationale. The core idea is that by first refining the pool through progressive filtering, coverage-based methods can focus on selecting diverse instances from a high-value candidate set rather than attempting to balance value and diversity over the entire unlabeled pool. Without this separation, a single strategy must typically simultaneously determine instance value and ensure batch diversity, which is a challenging multi-objective optimization problem. By operating on the refined pool $\mathcal{C}_R$, we can assume most of the candidates are valuable and focus on maximizing the coverage. This is both more efficient, as $|\mathcal{C}_R| \ll |\mathcal{U}_t|$, and more effective, since progressive filtering has already eliminated most uninformative instances.

Limitations. The candidate set $\mathcal{C}_R$ naturally reflects the collective emphasis of strategy set $\mathcal{S}$. When $\mathcal{S}$ contains more representativeness-based strategies, $\mathcal{C}_R$ will have

more representative instances. In such cases, incorporating uncertainty (as done in UHerd) helps to maintain an appropriate balance between exploration and exploitation. Conversely, when $\mathcal{S}$ contains more uncertainty-based strategies, the coverage optimization itself may provide the necessary representativeness.

4.3. Theoretical Analysis of Progressive Filtering

We now analyze the key properties of progressive filtering. We show that, with high probability, it (i) preserves value through union operations and (ii) exponentially reduces uninformative instances via iterative filtering. Furthermore, we show that it (iii) guarantees that the expected true value of the candidate pool does not decrease. Proofs and numerical examples can be found in Appendix A.

Value Preservation. The union operation in Eq. (1) ensures that instances deemed valuable by at least one strategy are likely to be preserved in the candidate pool.

Theorem 1 (Value Preservation Bound). *Let $\mathbf{x} \in \mathcal{C}_{r-1}$ be an instance and let $p_{m,r}(\mathbf{x})$ be the probability that strategy $s_m \in \mathcal{S}$ includes $\mathbf{x}$ in a selected batch during round r , given that $\mathbf{x}$ is present in the random subsample.*

The probability of $\mathbf{x}$ surviving round r , $P_r(\mathbf{x}) = \Pr(\mathbf{x} \in \mathcal{C}_r \mid \mathbf{x} \in \mathcal{C}_{r-1})$, is lower-bounded by the probability from the single strategy that is most likely to select $\mathbf{x}$:

$$P_r(\mathbf{x}) \geq 1 - \left(1 - \alpha \cdot \max_{m \in \{1, \dots, M\}} p_{m,r}(\mathbf{x})\right)^J.$$

Theorem 1 quantifies how progressive filtering avoids accidentally discarding valuable instances. An instance $\mathbf{x}$ valued by any strategy (high $\max_m p_{m,r}(\mathbf{x})$) has its survival probability lower-bounded by the J independent trials of the strategy it values most. This ensures that instances valued by any strategy are highly likely to remain in the pool, even with random subsampling. Importantly, the hyperparameters J and α offer a way to control this behavior.

Exponential Reduction of Uninformative Instances.

Progressive filtering leads to an exponential reduction of uninformative instances through successive refinement. We formalize the uninformativeness of an instance through a probabilistic model.

Definition 1 (Uninformative Instance). *An instance $\mathbf{x} \in \mathcal{U}_t$ is ϵ -uninformative if for all strategies $s_m \in \mathcal{S}$ and rounds r , the probability that $\mathbf{x}$ appears in a randomly selected batch, given $\mathbf{x}$ is present in the random subsample, is at most ϵ .*

Intuitively, uninformative instances are those that rank consistently low across all strategies and rounds and are not selected unless by random chance.

Theorem 2 (Exponential Reduction of Uninformative Instances). *Let $\mathbf{x}$ be an ϵ -uninformative instance. The probability that $\mathbf{x}$ survives R rounds of progressive filtering is*

bounded by:

$$\Pr(\mathbf{x} \in \mathcal{C}_R) \leq (1 - (1 - \alpha\epsilon)^{MJ})^R.$$

For $\alpha\epsilon \ll \frac{1}{MJ}$ and large R , this probability decreases exponentially as $\mathcal{O}((MJ\alpha\epsilon)^R)$.

Theorem 2 quantifies how the iterative process exponentially filters out uninformative instances. Since an uninformative instance is by definition unlikely to be selected by any strategy, it must survive the filtering process repeatedly across all R rounds to remain in the pool. Repeated filtering ensures that the survival probability for such an instance decreases exponentially with each round.

Concentration Toward a High-Value Pool. Beyond preserving valuable instances and filtering out uninformative ones, progressive filtering ensures that the value of the candidate pool does not decrease over rounds. We formalize this using the expected value of the instances in the pool.

Theorem 3 (Concentration of Expected Value). *Let $V(\mathbf{x})$ be the unknown true value of an instance $\mathbf{x}$, and let $\mathbb{E}[V|\mathcal{C}] = \frac{1}{|\mathcal{C}|} \sum_{\mathbf{x} \in \mathcal{C}} V(\mathbf{x})$ be the expected average value of pool $\mathcal{C}$. Assume that all strategies $s_m \in \mathcal{S}$ are more likely to select higher-value instances:*

$$V(\mathbf{x}) > V(\mathbf{y}) \implies p_{m,r}(\mathbf{x}) \geq p_{m,r}(\mathbf{y}) \quad \forall m, r$$

Then, the expected average value of the candidate pool is monotonically non-decreasing with each round of filtering:

$$\mathbb{E}[V \mid \mathcal{C}_R] \geq \mathbb{E}[V \mid \mathcal{C}_{R-1}] \geq \dots \geq \mathbb{E}[V \mid \mathcal{C}_0],$$

where $\mathcal{C}_0 = \mathcal{U}_t$.

Theorem 3 guarantees that the pool’s average value is monotonically non-decreasing, *i.e.*, it is guaranteed to improve or stay the same with each round. The key assumption, that AL strategies favor higher-value instances, is natural for any reasonable selection heuristic. This theorem provides the theoretical view for why progressive filtering concentrates the pool toward high-value instances. In each round, lower-value instances are filtered out at higher rates than higher-value ones, thereby shifting the distribution of the pool toward instances with greater true value. Together with the other theorems, this ensures efficient convergence to a high-quality candidate pool for the second stage.

5. Experimental Evaluation

We evaluate REFINE on image classification tasks, focusing primarily on standard benchmarks (*e.g.*, CIFAR-10). We also present a use case study on audio spectrogram classification in Appendix B. We start by establishing the experimental setup, then demonstrate that REFINE achieves superior performance against state-of-the-art strategies, followed by ablation studies analyzing key design choices. Our implementation can be found at [anonymous](#).¹

¹Our code will be made publicly available upon acceptance.

5.1. Experimental Setup

Datasets and AL Protocol. We evaluate REFINE on 6 image classification datasets. Following the AL protocol of [16], we run 20 cycles starting with a randomly sampled initial labeled pool of size b . The batch sizes per datasets were determined by ensuring that accuracy converges with random sampling. Our benchmark includes datasets with varying numbers of classes (from 10 to 200) to assess robustness across different complexity levels. As CIFAR-10 was used for determining R , J , and α , the *main results* report performance on the remaining datasets. Details are summarized in Tab. 1.

Table 1. Overview of datasets, showing number of instances, number of classes K and batch size b . CIFAR-10* was used for hyperparameter selection.

Dataset	# Instances	# Classes (K)	Batch Size (b)
CIFAR-10* [2009]	50,000	10	10
Dopanim [2024]	10,484	15	50
Snacks [2021]	4,838	20	20
CIFAR-100 [2009]	50,000	100	100
Food101 [2014]	75,750	101	100
Tiny ImageNet [2015]	100,000	200	200

Models. We employ 3 vision foundation models, each appended with a randomly initialized fully-connected classification head. We use DINOv2-ViT-S/14 [22], DINOv3-ViT-S/16 [30], and CLIP-ViT-B/16 [25]. After each batch selection, we fine-tune the classification head for 200 epochs while keeping the foundation model frozen. We use SGD with a batch size 64, learning rate 0.01, weight decay 10^{-4} , and cosine annealing scheduling. These values were chosen to ensure consistent convergence across all datasets.

Baseline Strategies. We compare REFINE against random sampling and 8 state-of-the-art AL strategies. Our baselines include: (i) the *uncertainty-based* strategies Margin [29] and BADGE [2], (ii) the *diversity-based* strategies TypiClust [12] and MaxHerding [3], and (iii) the *hybrid* approaches BAIT [1, 15], ALFAMix [23], DropQuery [10] and UHerding [4]. Finally, we include SelectAL [11], TCM [9], and TAILOR [31], recent ensemble AL methods that combine multiple AL strategies.

Evaluation Metrics. We assess AL performance using two metrics: (i) *relative learning curves*, showing accuracy gains (absolute percentage points) over random sampling (from $\mathcal{U}_t$) at each cycle, and (ii) *area under the learning curve* (AULC), measuring cumulative performance across all cycles. Absolute learning curves are provided in Appendix D. All results are averaged over 10 independent trials and standard errors are omitted for visibility. Additionally, we report *pairwise win rates* between strategies: the percentage of trials where strategy i achieves higher AULC than strategy j across all cycles.

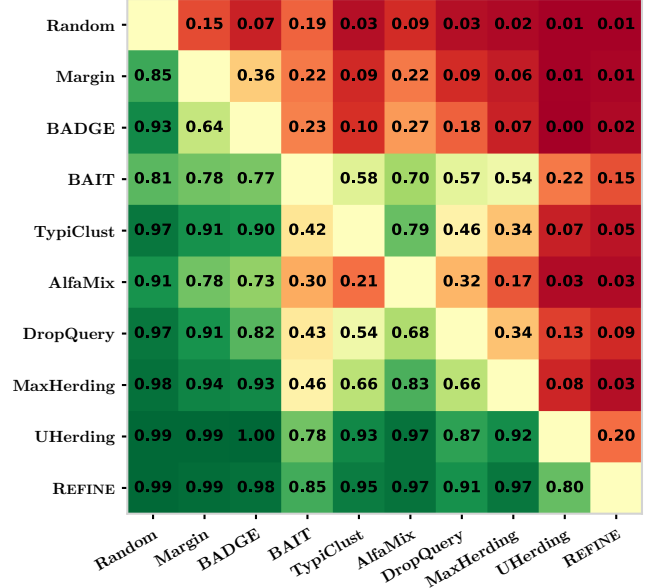


Figure 2. Pairwise comparison matrix averaged across 3 backbones, 5 datasets, and 10 trials. Element (i, j) corresponds to the proportion of total runs, where strategy i outperforms strategy j .

5.2. Main Results

Performance against Individual Strategies. Figure 2 averages the results across all backbones and datasets in a pairwise comparison matrix, where each element (i, j) reports the proportion of trials in which strategy i outperforms strategy j . Overall, across all backbones and datasets, REFINE emerges as the most competitive strategy, achieving the highest win rates and outperforming every other strategy. For example, compared to other top-performing AL strategies like BAIT and UHerding, REFINE wins in 85% and 80% of all cases, respectively.

Additionally, Fig. 3 shows an exemplary subset of relative learning curves on the four most complex datasets, focusing on DINOv3 as the newest model and CLIP for its differing training process (*i.e.*, student-teacher vs. contrastive learning). The complete set of learning curves (relative and absolute) can be found in Appendix D. We observe that REFINE consistently outperforms all strategies on CIFAR-100, Food101, and Tiny ImageNet across all backbones, demonstrating its robustness across different datasets and models.

Importantly, we also want to highlight a limitation of progressive filtering: the performance can degrade when the majority of strategies in its ensemble perform poorly. This can be seen on the Dopanim dataset using the CLIP backbone, where progressive filtering does not yield the best learning curve. Specifically, at around 350 labels, the accuracy of REFINE decreases slightly, which occurs as three AL strategies begin to perform equal to or worse than random sampling. While this is a general problem of en-

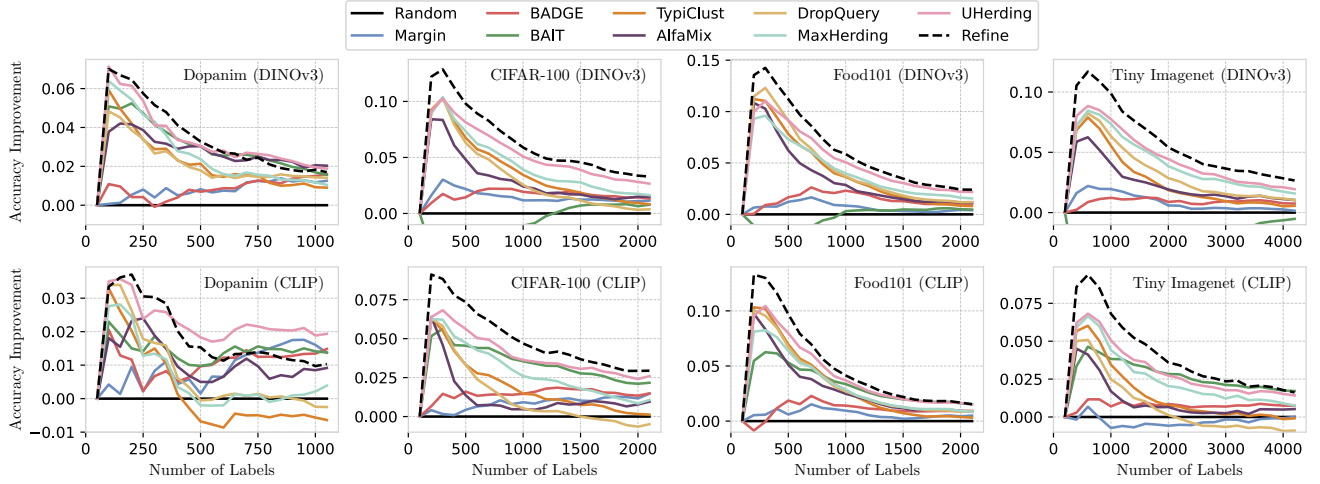


Figure 3. Relative accuracy learning curves for REFINE and baseline strategies across multiple backbones and datasets.

semble methods, it could potentially be mitigated through weighting of ensemble components, which we identify as a promising direction for future research.

Performance against Ensemble Methods. Figure 4 extends the pairwise comparison to ensemble AL methods. Here, we see that REFINE clearly outperforms all other methods, achieving a 100% win rate against both SelectAL and TAILOR, and 98% against TCM. These results indicate that REFINE consistently combines AL strategies more effectively than competitors. The complete set of associated learning curves can be found in Appendix D.

Random		0.01	0.10	0.15	0.01
TCM	0.99		0.80	0.97	0.02
SelectAL	0.90	0.20		0.71	0.00
TAILOR	0.85	0.03	0.29		0.00
REFINE	0.99	0.98	1.00	1.00	
	Random	TCM	SelectAL	TAILOR	REFINE

Figure 4. Pairwise comparison matrix of ensemble AL methods averaged across 3 backbones, 5 datasets, and 10 trials.

5.3. Ablation Studies

We now investigate the contribution of REFINE’s key components through four research questions, varying both datasets and backbones to ensure our findings are not specific to a particular experimental configuration. Additional results are provided in Appendix C.

Does progressive filtering yield a higher-quality candidate pool? To isolate the contribution of progressive filtering from any specific AL strategy, we first apply filtering

Table 2. Effect of filtering rounds R , sampling ratio α , and batches per strategy J on pool quality. We report AULC improvement [%] of random sampling when applied to filtered vs. unfiltered pools.

(a) Filtering rounds R			(b) Sampling ratio α			(c) # Batches J		
R	CIFAR-10	Snacks	α	CIFAR-10	Snacks	J	CIFAR-10	Snacks
1	3.02	6.91	0.1	3.11	7.37	1	3.57	7.88
2	3.31	6.75	0.2	3.33	6.32	5	3.71	7.79
3	3.72	7.22	0.3	3.89	7.60	10	3.79	7.22
5	3.71	7.79	0.5	3.62	8.16	12	3.89	8.05
7	3.81	8.10	0.7	3.97	7.53	15	3.85	7.30
9	3.78	8.43	0.9	3.70	8.40	20	3.93	7.69

with $R \in \{1, 2, 3, 5, 7, 9\}$ and perform *random sampling* from the filtered pool $\mathcal{C}_R$. We fix $\alpha = 0.4$ and $J = 5$ throughout these experiments. Table 2a shows that progressive filtering consistently improves pool quality across CIFAR-10 and Snacks. We observe better accuracies compared to random sampling from $\mathcal{U}_t$, with performance increasing as more filtering rounds are applied. This can saturate with higher rounds, *e.g.*, for CIFAR-10 and $R = 9$, the improvement plateaus, indicating diminishing returns from additional filtering rounds. However, overall, increasing depth R generally yields a higher-quality candidate pool.

How do α and J affect performance of progressive filtering? We again isolate the contribution of progressive filtering from any specific selection strategy and examine α and J . Table 2b shows that random sampling with progressive filtering ($R = 5$ and $J = 5$) consistently outperforms random sampling from $\mathcal{U}_t$ across all values of α . Values that are too low ($\alpha \leq 0.2$) can slightly reduce filtering effectiveness, while higher values ($\alpha \in [0.3, 0.9]$) perform well for both datasets. Importantly, since α determines the pool size from which strategy s_m samples, lower values are preferable as they reduce computational and memory cost. Simi-

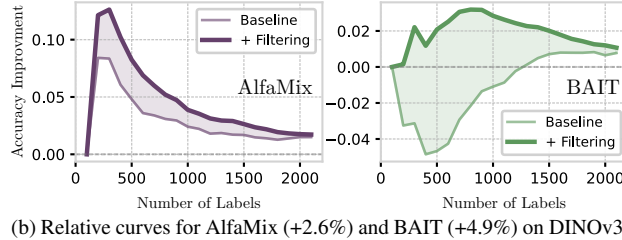
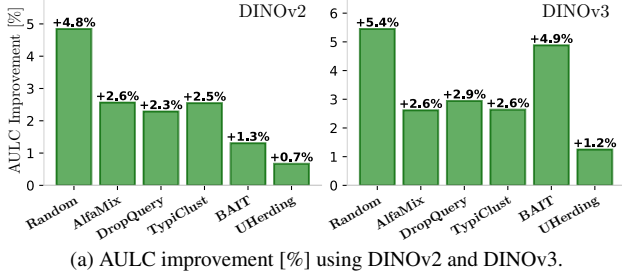


Figure 5. Benefits of progressive filtering (selecting from $\mathcal{C}_R$ vs. $\mathcal{U}_t$) on CIFAR-100 across different strategies.

larly, Tab. 2c demonstrates that progressive filtering ($R = 5$ and $\alpha = 0.4$) generally improves performance when increasing J . Accuracy improves as more batches are selected per strategy, but this benefit can diminish for larger values.

How beneficial is progressive filtering as a preprocessing step for AL strategies? We now examine how progressive filtering benefits different AL selection strategies. In REFINE, we employ coverage-based selection on the candidate pool $\mathcal{C}_R$ to efficiently cover valuable instances. Here, we isolate this second stage by first applying progressive filtering with fixed parameters ($R = 5$, $\alpha = 0.4$, $J = 5$), then comparing selection strategies from $\mathcal{S}$ when applied to the filtered pool $\mathcal{C}_R$ versus the unfiltered pool $\mathcal{U}_t$. Figure 5a shows AULC improvements for random sampling and five state-of-the-art selection strategies on CIFAR-100 using DINOv2 and DINOv3. Progressive filtering consistently improves all strategies for both backbones, with gains ranging from +0.7% to +5.4%, demonstrating that any selection strategy can benefit from our preprocessing step. Notably, weaker AL strategies benefit most from progressive filtering, while well-performing strategies like UHerding show smaller gains. Figure 5b shows exemplary learning curves for AlfaMix and BAIT on CIFAR-100 using DINOv3. While progressive filtering improves both strategies, the effect is critical for BAIT. Without filtering, BAIT (Baseline) performs even worse than random sampling. However, incorporating filtering improves its performance, showing the benefit of an ensemble of AL strategies.

Does progressive filtering successfully combine complementary notions of value? To investigate whether progressive filtering properly captures different notions of value in the final pool $\mathcal{C}_R$, we perform filtering with only

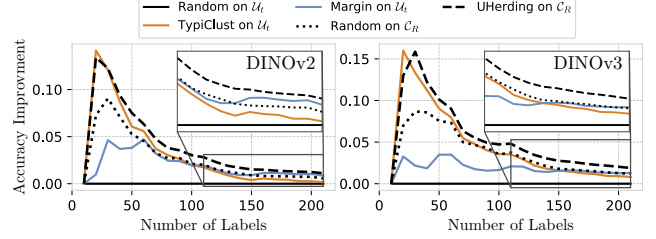


Figure 6. Relative learning curves on CIFAR-10 for strategies on $\mathcal{U}_t$ and $\mathcal{C}_R$ with filtering restricted to Margin and TypiClust.

two *complementary* selection strategies: Margin, focusing on uncertain instances, and TypiClust, focusing on representative instances. If filtering captures these notions correctly, this should be evident when randomly sampling from $\mathcal{C}_R$. We fix hyperparameters to the values used in our main experiments ($R = 5$, $\alpha = 0.4$, $J = 10$). Figure 6 demonstrates that the final pool $\mathcal{C}_R$ correctly captures both notions of value. Initially, TypiClust yields higher performances due to its focus on representative instances, and random sampling from $\mathcal{C}_R$ effectively profits from this. As TypiClust’s effectiveness diminishes in later cycles, random sampling from $\mathcal{C}_R$ benefits from uncertain instances captured by Margin. This behavior is consistent across both backbones. Moreover, coverage-based selection via UHerding on $\mathcal{C}_R$ further improves performance, demonstrating its ability to capture these diverse notions of value.

6. Conclusion

We introduced REFINE, an ensemble AL method that effectively combines different notions of value (*e.g.*, uncertainty and representativeness) to achieve state-of-the-art performance. Our core contribution, progressive filtering, iteratively refines the unlabeled pool via the union of batches from multiple AL strategies. This process cleans the unlabeled pool by discarding uninformative instances, yielding a high-value candidate pool for downstream selection. A subsequent coverage-based selection then effectively yields a diverse batch from this refined, high-quality pool.

Limitations. The effectiveness of REFINE depends on the quality and composition of its strategy ensemble. If these strategies are highly correlated (*e.g.*, all share a focus on diversity), the filtered pool will inherit this bias, diminishing the benefits of the ensemble. Moreover, the union can be sensitive if there are many poorly-designed strategies. As these will repeatedly select uninformative instances, this may slightly degrade the value of the candidate pool. We therefore recommend building $\mathcal{S}$ from a small set of well-established, complementary AL strategies rather than arbitrarily including numerous ones.

References

- [1] Jordan Ash, Surbhi Goel, Akshay Krishnamurthy, and Sham Kakade. Gone fishing: Neural active learning with fisher embeddings. In *NeurIPS*, 2021. 1, 2, 6
- [2] Jordan T. Ash, Chicheng Zhang, Akshay Krishnamurthy, John Langford, and Alekh Agarwal. Deep Batch Active Learning by Diverse, Uncertain Gradient Lower Bounds. In *ICLR*, 2020. 1, 2, 6
- [3] Wonho Bae, Junhyug Noh, and Danica J Sutherland. Generalized coverage for more robust low-budget active learning. In *ECCV*, 2024. 2, 4, 6
- [4] Wonho Bae, Danica J. Sutherland, and Gabriel L. Oliveira. Uncertainty herding: One active learning method for all label budgets. In *ICLR*, 2025. 1, 2, 4, 6
- [5] Yoram Baram, Ran El Yaniv, and Kobi Luz. Online choice of active learning algorithms. *JMLR*, 5:255–291, 2004. 2
- [6] Lukas Bossard, Matthieu Guillaumin, and Luc Van Gool. Food-101 – mining discriminative components with random forests. In *ECCV*, 2014. 6
- [7] Wenxi Chen, Yuzhe Liang, Ziyang Ma, Zhisheng Zheng, and Xie Chen. Eat: Self-supervised pre-training with efficient audio transformer. In *IJCAI*, 2024. 3
- [8] Pinar Donmez, Jaime G Carbonell, and Paul N Bennett. Dual strategy active learning. In *ECML*, 2007. 2
- [9] Paul Doucet, Benjamin Estermann, Till Aczel, and Roger Wattenhofer. Bridging diversity and uncertainty in active learning with self-supervised pre-training. In *PML4LRS @ ICLR*, 2024. 2, 6
- [10] Sanket Rajan Gupte, Josiah Aklilu, Jeffrey J. Nirschl, and Serena Yeung-Levy. Revisiting Active Learning in the Era of Vision Foundation Models. *TMLR*, 2024. 2, 6
- [11] Guy Hacohen and Daphna Weinshall. How to select which active learning strategy is best suited for your specific problem and budget. In *NeurIPS*, 2023. 2, 6
- [12] Guy Hacohen, Avihu Dekel, and Daphna Weinshall. Active learning on a budget: Opposite strategies suit high and low budgets. In *ICML*, 2022. 2, 4, 6
- [13] Marek Herde, Denis Huseljic, Lukas Rauch, and Bernhard Sick. dopanim: A dataset of doppelganger animals with noisy annotations from multiple humans. In *NeurIPS*, 2024. 6
- [14] Wei-Ning Hsu and Hsuan-Tien Lin. Active learning by learning. In *AAAI*, 2015. 2
- [15] Denis Huseljic, Paul Hahn, Marek Herde, Lukas Rauch, and Bernhard Sick. Fast fishing: Approximating bait for efficient and scalable deep active image classification. In *ECML*, 2024. 6
- [16] Denis Huseljic, Marek Herde, Lukas Rauch, Paul Hahn, Zhixin Huang, Daniel Kottke, Stephan Vogt, and Bernhard Sick. Efficient bayesian updates for deep active learning via laplace approximations. In *ECML*, 2025. 6
- [17] Alex Krizhevsky. Learning Multiple Layers of Features from Tiny Images. Master’s thesis, University of Toronto, 2009. 6
- [18] Ya Le and Xuan Yang. Tiny ImageNet Visual Recognition Challenge. *CS231N Course Project Report*, 2015. 6
- [19] Carsten Lüth, Till Bungert, Lukas Klein, and Paul Jaeger. Navigating the pitfalls of active learning evaluation: A systematic framework for meaningful performance assessment. *NeurIPS*, 36, 2024. 1
- [20] Matthijs. Snacks dataset. <https://huggingface.co/datasets/Matthijs/snacks>, 2021. Accessed: 2024-05-20. 6
- [21] Prateek Munjal, Nasir Hayat, Munawar Hayat, Jamshid Sourati, and Shadab Khan. Towards robust and reproducible active learning using neural networks. In *CVPR*, 2022. 1
- [22] Maxime Oquab, Timothée Darcet, Théo Moutakanni, Huy Vo, Marc Szafraniec, Vasil Khalidov, Pierre Fernandez, Daniel Haziza, Francisco Massa, Alaaeldin El-Nouby, et al. Dinov2: Learning robust visual features without supervision. *TMLR*, 2024. 6
- [23] Amin Parvaneh, Ehsan Abbasnejad, Damien Teney, Gholamreza Reza Haffari, Anton van den Hengel, and Javen Qinfeng Shi. Active learning by feature mixing. In *CVPR*, 2022. 2, 6
- [24] Karol J. Piczak. ESC: Dataset for environmental sound classification. In *ACM*, 2015. 3
- [25] Alec Radford, Jong Wook Kim, Chris Hallacy, Aditya Ramesh, Gabriel Goh, Sandhini Agarwal, Girish Sastry, Amanda Askell, Pamela Mishkin, Jack Clark, Gretchen Krueger, and Ilya Sutskever. Learning transferable visual models from natural language supervision. In *ICML*, 2021. 6
- [26] Lukas Rauch, Matthias Aßenmacher, Denis Huseljic, Moritz Wirth, Bernd Bischl, and Bernhard Sick. Activeglac: A benchmark for deep active learning with transformers. In *ECML*, 2023. 1
- [27] Lukas Rauch, René Heinrich, Houtan Ghaffari, Lukas Miklautz, Ilyass Moummad, Bernhard Sick, and Christoph Scholz. Unmute the patch tokens: Rethinking probing in multi-label audio classification. *arXiv preprint arXiv:2509.24901*, 2025. 3
- [28] Ozan Sener and Silvio Savarese. Active Learning for Convolutional Neural Networks: A Core-Set Approach. In *ICLR*, 2018. 4, 3
- [29] Burr Settles. Active Learning Literature Survey. Technical report, University of Wisconsin, Department of Computer Science, 2009. 1, 2, 6
- [30] Oriane Siméoni, Huy V Vo, Maximilian Seitzer, Federico Baldassarre, Maxime Oquab, Cijo Jose, Vasil Khalidov, Marc Szafraniec, Seungeun Yi, Michaël Ramamonjisoa, et al. Dinov3. *arXiv preprint arXiv:2508.10104*, 2025. 1, 6
- [31] Jifan Zhang, Shuai Shao, Saurabh Verma, and Robert Nowak. Algorithm selection for deep active learning with imbalanced datasets. In *NeurIPS*, 2023. 2, 6

Cleaning the Pool: Progressive Filtering of Unlabeled Pools in Deep Active Learning

Supplementary Material

A. Proofs of Theorems

Here, we provide detailed proofs for the theorems from the main paper. Furthermore, to complement the theoretical analysis of Theorems 1 and 2, we also include numerical examples illustrating the probability of preserving valuable instances and filtering out uninformative ones.

Proof of Theorem 1. Let $P_r(\mathbf{x}) = \Pr(\mathbf{x} \in \mathcal{C}_r \mid \mathbf{x} \in \mathcal{C}_{r-1})$ be the probability of instance $\mathbf{x}$ surviving round r . Given $\mathbf{x} \in \mathcal{C}_{r-1}$, let $S_{m,r}$ be the *event* that AL strategy s_m selects $\mathbf{x}$ in round r at least once in its J trials. The total survival event is the union of these individual events, *i.e.*, $\mathbf{x}$ survives if at least one strategy selects it:

$$P_r(\mathbf{x}) = \Pr\left(\bigcup_{m=1}^M S_{m,r}\right).$$

We first derive $\Pr(S_{m,r})$ and then use this to find the lower bound for the union.

1. For a single draw $j \in \{1, \dots, J\}$ by strategy s_m , the event of selecting $\mathbf{x}$ requires two *independent conditions* to be met:
 - (a) Instance $\mathbf{x}$ must be present in the random subsample $\mathcal{U}_{\text{sample}}$. The probability of this is

$$\Pr(\mathbf{x} \in \mathcal{U}_{\text{sample}}) = \alpha.$$

- (b) Strategy s_m must select $\mathbf{x}$ from that subsample. By definition, this probability is $p_{m,r}(\mathbf{x})$.

2. The joint probability of $\mathbf{x}$ being selected in a single draw j is the product of these probabilities:

$$\Pr(\text{select } \mathbf{x} \text{ in draw } j) = \alpha \cdot p_{m,r}(\mathbf{x}).$$

3. The probability of $\mathbf{x}$ *not being selected* in this single draw is $1 - \alpha \cdot p_{m,r}(\mathbf{x})$.
4. As per Algorithm 1, strategy s_m makes J independent draws (each with a new subsample). The probability that $\mathbf{x}$ is *never selected* by s_m across all J draws is the product of their individual probabilities:

$$\Pr(\bar{S}_{m,r}) = \prod_{j=1}^J (1 - \alpha \cdot p_{m,r}(\mathbf{x})) = (1 - \alpha \cdot p_{m,r}(\mathbf{x}))^J.$$

5. Therefore, the probability that $\mathbf{x}$ is selected *at least once* by strategy s_m is the complement:

$$\Pr(S_{m,r}) = 1 - \Pr(\bar{S}_{m,r}) = 1 - (1 - \alpha \cdot p_{m,r}(\mathbf{x}))^J.$$

Having found $\Pr(S_{m,r})$, we now derive the lower bound.

1. The probability of a union of events is always greater than or equal to the probability of any single event in that union. Therefore, for any strategy s_i :

$$P_r(\mathbf{x}) = \Pr\left(\bigcup_{m=1}^M S_{m,r}\right) \geq \Pr(S_{i,r}).$$

2. Since this holds for all i , it must also hold for the maximum probability among them:

$$P_r(\mathbf{x}) \geq \max_{m \in \{1, \dots, M\}} \Pr(S_{m,r}).$$

3. Substituting the expression for $\Pr(S_{m,r})$ from Step 5:

$$P_r(\mathbf{x}) \geq \max_{m \in \{1, \dots, M\}} [1 - (1 - \alpha \cdot p_{m,r}(\mathbf{x}))^J].$$

4. The function $f(p) = 1 - (1 - \alpha p)^J$ is monotonically increasing with p . Therefore, the maximum value of the function is achieved when p is at its maximum.

Let $p_{\max} = \max_{m \in \{1, \dots, M\}} p_{m,r}(\mathbf{x})$. Then:

$$\max_m [1 - (1 - \alpha \cdot p_{m,r}(\mathbf{x}))^J] = 1 - (1 - \alpha \cdot p_{\max})^J.$$

5. This proves the final bound:

$$P_r(\mathbf{x}) \geq 1 - \left(1 - \alpha \cdot \max_{m \in \{1, \dots, M\}} p_{m,r}(\mathbf{x})\right)^J. \quad \blacksquare$$

Numerical Example. We provide a numerical example to illustrate the probability of preserving a valuable instance. We consider $\alpha = 0.4$, $J = 10$, and an instance $\mathbf{x}$ that receives a high score from Margin ($p_{\text{margin}} = 0.9$) but a low score from TypiClust ($p_{\text{typiclust}} = 0.2$). We assume an ensemble of two strategies, where the instance's true value is correlated with the uncertainty heuristic, which is accounted for by the Margin. Since the bound uses the maximum score, we have $p_{\max} = 0.9$. This yields $P_r(\mathbf{x}) \geq 1 - (1 - 0.4 \cdot 0.9)^{10} = 1 - (0.64)^{10} \approx 0.9885$, demonstrating that the instance has a high lower-bound probability because at least one strategy strongly prioritizes it.

Proof of Theorem 2. Let $\mathbf{x}$ be an ϵ -uninformative instance and let $P_r(\mathbf{x}) = \Pr(\mathbf{x} \in \mathcal{C}_r \mid \mathbf{x} \in \mathcal{C}_{r-1})$ be the probability that an ϵ -uninformative instance $\mathbf{x}$ survives round r . We first derive an upper bound for $P_r(\mathbf{x})$ and then use this to find the bound for surviving all R rounds.

1. For a single draw $j \in \{1, \dots, J\}$ by strategy s_m , the event of selecting $\mathbf{x}$ requires *two conditions* to be met:

- (a) Instance $\mathbf{x}$ must be present in the random subsample $\mathcal{U}_{\text{sample}}$. The probability of this is

$$\Pr(\mathbf{x} \in \mathcal{U}_{\text{sample}}) = \alpha.$$

- (b) Strategy s_m must select $\mathbf{x}$ from that subsample. By Definition 1, this probability is bounded: $p_{m,r}(\mathbf{x}) \leq \epsilon$.

2. The joint probability of $\mathbf{x}$ being selected in a single draw j by strategy s_m is bounded:

$$\Pr(\text{select } \mathbf{x} \text{ in draw } (m, j)) = \alpha \cdot p_{m,r}(\mathbf{x}) \leq \alpha\epsilon.$$

3. The probability of $\mathbf{x}$ *not* being selected in this single draw is therefore lower-bounded:

$$\Pr(\text{not select } \mathbf{x} \text{ in draw } (m, j)) \geq 1 - \alpha\epsilon.$$

4. As per Algorithm 1, we make a total of $M \cdot J$ independent draws in round r . The event that $\mathbf{x}$ is *never selected* requires it to be missed by all $M \cdot J$ draws. Using the lower bound from the previous step, we get a lower bound for $\mathbf{x}$ not surviving the round:

$$\Pr(\mathbf{x} \notin \mathcal{C}_r \mid \mathbf{x} \in \mathcal{C}_{r-1}) \geq \prod_{m=1}^M \prod_{j=1}^J (1 - \alpha\epsilon) = (1 - \alpha\epsilon)^{MJ}.$$

5. Consequently, the event that $\mathbf{x}$ is selected at least once is bounded by:

$$\begin{aligned} P_r(\mathbf{x}) &= \Pr(\mathbf{x} \in \mathcal{C}_r \mid \mathbf{x} \in \mathcal{C}_{r-1}) \\ &= 1 - \Pr(\mathbf{x} \notin \mathcal{C}_r \mid \mathbf{x} \in \mathcal{C}_{r-1}) \\ &\leq 1 - (1 - \alpha\epsilon)^{MJ}. \end{aligned}$$

Having found the upper bound for $P_r(\mathbf{x})$, we now derive the bound for surviving all R rounds.

1. For $\mathbf{x}$ to survive all R rounds, it must survive each round. The total survival probability is the product of the per-round survival probabilities:

$$\Pr(\mathbf{x} \in \mathcal{C}_R) = \prod_{r=1}^R \Pr(\mathbf{x} \in \mathcal{C}_r \mid \mathbf{x} \in \mathcal{C}_{r-1}) = \prod_{r=1}^R P_r(\mathbf{x}).$$

2. Since, by definition, instances $\mathbf{x}$ remain ϵ -uninformative across rounds (for all r):

$$\begin{aligned} \Pr(\mathbf{x} \in \mathcal{C}_R) &\leq \prod_{r=1}^R (1 - (1 - \alpha\epsilon)^{MJ}) \\ &= (1 - (1 - \alpha\epsilon)^{MJ})^R, \end{aligned}$$

proving the first part of Theorem 2.

3. For the approximation, we assume $\alpha\epsilon \ll \frac{1}{MJ}$. Using the approximation $(1 - x)^n \approx 1 - nx$ for small x :

$$(1 - \alpha\epsilon)^{MJ} \approx 1 - MJ\alpha\epsilon.$$

4. Substituting this approximation back into the bound, we obtain that

$$\Pr(\mathbf{x} \in \mathcal{C}_R) \lesssim (1 - (1 - MJ\alpha\epsilon))^R = (MJ\alpha\epsilon)^R,$$

showing that the probability decreases exponentially with R . ■

Numerical Example. We provide a numerical example to illustrate the survival probability of an uninformative instance. We set $R = 5$, $M = 3$, $J = 10$, and $\alpha = 0.4$. Furthermore, we consider an ϵ -uninformative instance $\mathbf{x}$ with $\epsilon = 0.05$. The probability that this instance survives a single round is

$$\begin{aligned} P_r(\mathbf{x}) &\leq 1 - (1 - \alpha\epsilon)^{MJ} \\ &= 1 - (1 - 0.4 \cdot 0.05)^{30} \\ &= 1 - (0.98)^{30} \approx 0.4545. \end{aligned}$$

Extending this over all $R = 5$ rounds, the bound becomes $\Pr(\mathbf{x} \in \mathcal{C}_5) \leq (1 - (0.98)^{30})^5 \approx 0.0194$. After 5 rounds, the probability of this uninformative instance remaining in the candidate pool is bounded below 2%.

Proof of Theorem 3. Let $V(\mathbf{x})$ be the unknown true value of instance $\mathbf{x}$, and let $\mathbb{E}[V|\mathcal{C}] = \frac{1}{|\mathcal{C}|} \sum_{\mathbf{x} \in \mathcal{C}} V(\mathbf{x})$ be the expected average value of pool $\mathcal{C}$. We show that the expected value is non-decreasing in any single round r .

- Let $\mathbb{E}[V \mid \mathcal{C}_{r-1}] = \frac{1}{|\mathcal{C}_{r-1}|} \sum_{\mathbf{x} \in \mathcal{C}_{r-1}} V(\mathbf{x})$ be the average value of the pool at the start of the round.
- The pool $\mathcal{C}_r$ is formed by sampling from $\mathcal{C}_{r-1}$. As shown in the proof of Theorem 1, the probability that an instance $\mathbf{x} \in \mathcal{C}_{r-1}$ survives to $\mathcal{C}_r$ is

$$P_r(\mathbf{x}) = \Pr(\mathbf{x} \in \mathcal{C}_r \mid \mathbf{x} \in \mathcal{C}_{r-1}).$$

- The expected value of the new pool $\mathcal{C}_r$ is the weighted average of the values from the previous pool, where the weights are these survival probabilities²:

$$\mathbb{E}[V \mid \mathcal{C}_r] = \frac{\sum_{\mathbf{x} \in \mathcal{C}_{r-1}} V(\mathbf{x}) P_r(\mathbf{x})}{\sum_{\mathbf{x} \in \mathcal{C}_{r-1}} P_r(\mathbf{x})}.$$

- From the theorem's assumption, an instance's value $V(\mathbf{x})$ is positively correlated with its selection probability $p_{m,r}(\mathbf{x})$ for all strategies:

$$V(\mathbf{x}) > V(\mathbf{y}) \implies p_{m,r}(\mathbf{x}) \geq p_{m,r}(\mathbf{y}) \quad \forall m, r.$$

- From this and the fact that $P_r(\mathbf{x})$ is a monotonically increasing function of each $p_{m,r}(\mathbf{x})$, it follows that $V(\mathbf{x})$ is *also positively correlated* with $P_r(\mathbf{x})$:

$$V(\mathbf{x}) > V(\mathbf{y}) \implies P_r(\mathbf{x}) \geq P_r(\mathbf{y}).$$

²Since we iterate over the specific instances fixed in the previous pool $\mathcal{C}_{r-1}$, the value $V(\mathbf{x})$ is treated as a constant.

6. Because the values $V(\mathbf{x})$ and the weights $P_r(\mathbf{x})$ are positively correlated, *Chebyshev's sum inequality* shows that the weighted average must be greater than or equal to the unweighted average:

$$\underbrace{\frac{\sum_{\mathbf{x} \in \mathcal{C}_{r-1}} V(\mathbf{x}) P_r(\mathbf{x})}{\sum_{\mathbf{x} \in \mathcal{C}_{r-1}} P_r(\mathbf{x})}}_{\mathbb{E}[V | \mathcal{C}_r]} \geq \underbrace{\frac{\sum_{\mathbf{x} \in \mathcal{C}_{r-1}} V(\mathbf{x})}{|\mathcal{C}_{r-1}|}}_{\mathbb{E}[V | \mathcal{C}_{r-1}]}$$

7. Since $\mathbb{E}[V | \mathcal{C}_r] \geq \mathbb{E}[V | \mathcal{C}_{r-1}]$ holds for any round r , applying this result iteratively proves:

$$\mathbb{E}[V | \mathcal{C}_R] \geq \mathbb{E}[V | \mathcal{C}_{R-1}] \geq \dots \geq \mathbb{E}[V | \mathcal{C}_0]. \quad \blacksquare$$

B. Use Case: Audio Spectrogram Classification

This use case demonstrates the practical utility of progressive filtering as a preprocessing step in AL. Specifically, we demonstrate that practitioners can achieve substantial accuracy gains, even transforming a previously underperforming strategy into a viable one. They can accomplish this without discarding their existing, trusted AL workflows by employing progressive filtering as a workflow-agnostic method to refine the unlabeled pool. This approach enhances any AL workflow with minimal engineering effort and risk. An additional benefit can also be computational feasibility. By reducing the unlabeled pool size, progressive filtering enables the practical application of complex or computationally expensive AL strategies that would otherwise be prohibitively costly on large (unfiltered) unlabeled pools.

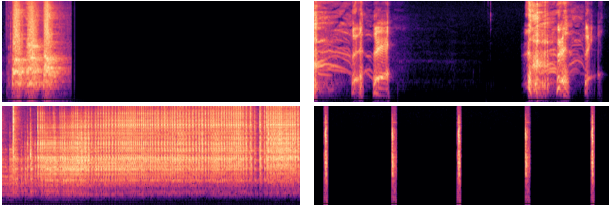


Figure 7. Example spectrograms from ESC50.

To illustrate this, we conducted a use case study on the ESC50 audio dataset [24]. As ESC50 does not have a dedicated test split, we randomly sampled 20% of the data for evaluation. We transformed audio segments into spectrograms following [27] (*cf.* Fig. 7) and employed the EAT base backbone to extract features [7]. Progressive Filtering was applied exactly as described in the main paper, using the same hyperparameters and ensemble $\mathcal{S}$. This use case simulates a practitioner whose trusted AL strategy is Coreset [28], which is highly cited and well-known. Critically, Coreset is not part of the ensemble $\mathcal{S}$ used for filtering, which is a scenario where progressive filtering enhances an external strategy.

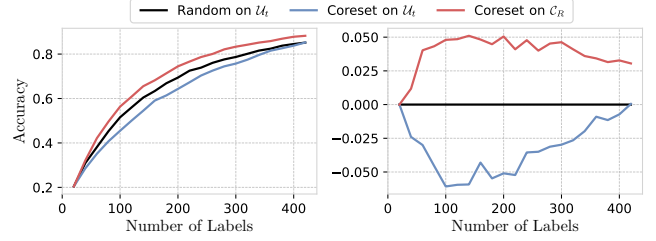


Figure 8. Absolute and relative learning curves comparing random sampling on $\mathcal{U}_t$, Coreset on $\mathcal{U}_t$, Coreset on $\mathcal{C}_R$ on ESC50.

As shown in Figure 8, applying the original Coreset strategy to the entire unlabeled pool $\mathcal{U}_t$ yields accuracy considerably below the random sampling baseline. In this instance, the standard strategy failed, meaning the practitioner would have obtained better results using simple random sampling. In contrast, by applying progressive filtering as a preprocessing step, we observe an accuracy improvement of up to 10%. This demonstrates that filtering can even effectively “fix” a failing strategy, allowing it to outperform baselines (*e.g.*, random) without requiring practitioners to abandon their preferred workflow.

C. Ablation Studies: Additional Ablations

In Fig. 5a, we reported improvements in AULC [%], which showed that all strategies applied to the progressively filtered pool $\mathcal{C}_R$ lead to performance improvement. For these strategies, Fig. 9 presents the corresponding accuracy learning curves. The figure compares strategies applied to the filtered pool versus the unfiltered pool.

Figure 9 demonstrates that progressive filtering acts as a powerful preprocessing step, consistently improving the accuracies of individual AL strategies. This behavior is consistent across both DINOv2 and DINOv3 backbones. While the impact is most critical for poorly performing AL strategies, progressive filtering still enhances strong performers such as UHerding, even if the improvement is less pronounced. Notably, filtering significantly improves even random sampling, demonstrating its ability to clean uninformative instances from the unlabeled pool. Furthermore, progressive filtering can act as a crucial stabilizer for failing AL strategies. For example, while standard BAIT performs worse than random sampling on DINOv3, applying filtering fixes this failure and substantially improves the strategy’s effectiveness.

D. Main Results: All Learning Curves

This section presents the complete set of learning curves used to derive Figs. 2 and 4, which were omitted from the main paper due to space limitations. For individual AL strategies, Fig. 10 displays the absolute learning curves (accuracy), while Fig. 11 shows the relative learning curves

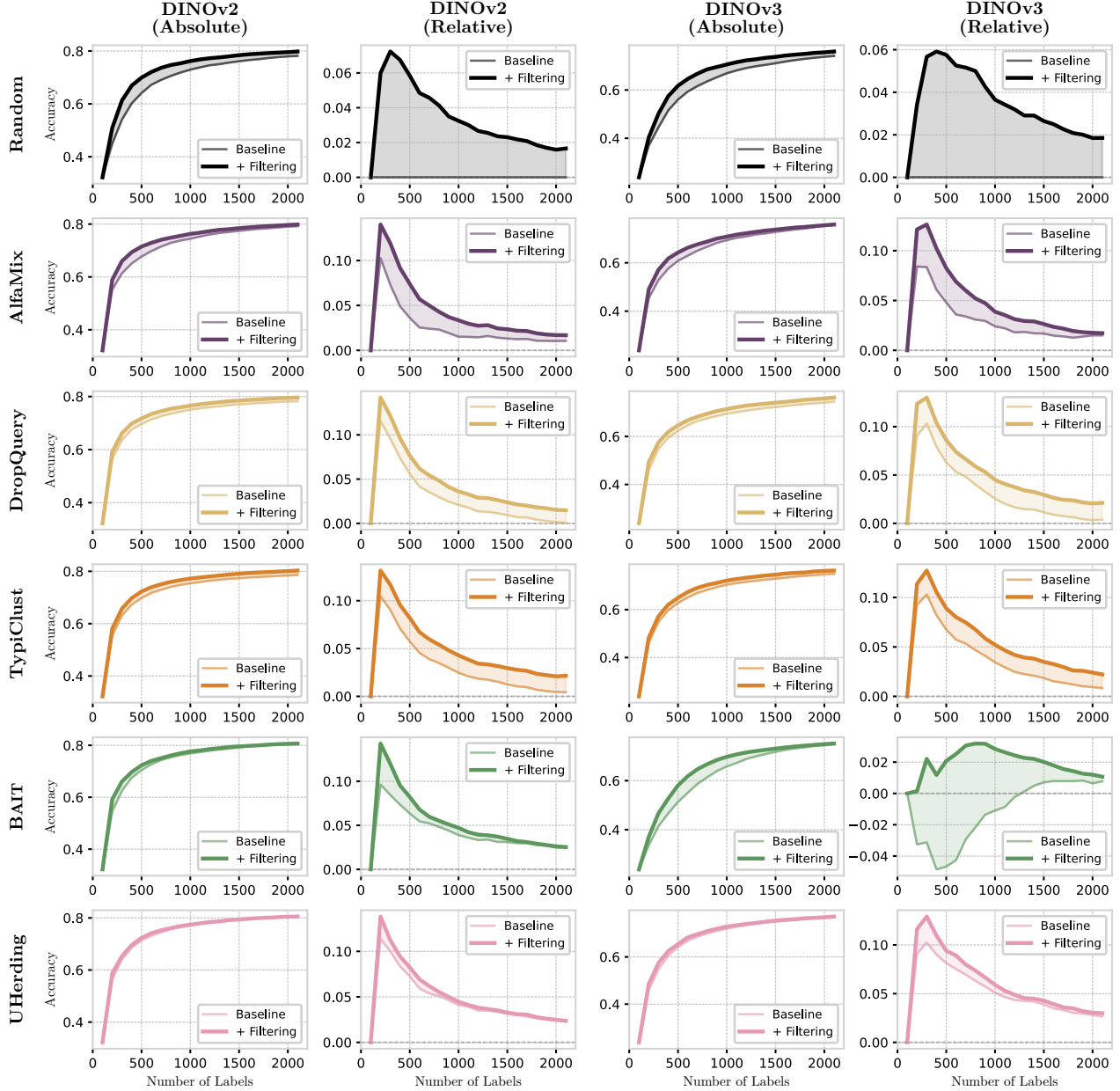


Figure 9. Relative and absolute learning curves on CIFAR-100 comparing selection strategies applied to the filtered pool (bold lines) versus the unfiltered pool (normal lines) using DINOv2 and DINOv3 backbones.

(accuracy relative to random sampling from $\mathcal{U}_t$). Similarly, for ensemble AL methods, Figs. 12 and 13 present the absolute and relative learning curves, respectively. These results provide a comprehensive overview supporting the conclusions drawn in the main text.

Figures 10 and 11 demonstrate that REFINE outperforms all other individual AL strategies across nearly every dataset and backbone, including DINOv2, CLIP, and DINOv3. It excels in both low- and high-budget scenarios, consistently delivering high accuracies throughout the entire AL pro-

cess. Furthermore, REFINE proves to be robust, maintaining a strong lead across diverse datasets while remaining invariant to the choice of the backbone. Similarly, Figures 12 and 13 demonstrate that REFINE outperforms all ensemble AL methods. These approaches generally struggle to effectively combine multiple strategies, often failing to exceed the performance of individual AL strategies. With REFINE, we propose the first ensemble AL method capable of effectively combining multiple AL strategies into a unified framework that consistently surpasses individual baselines.

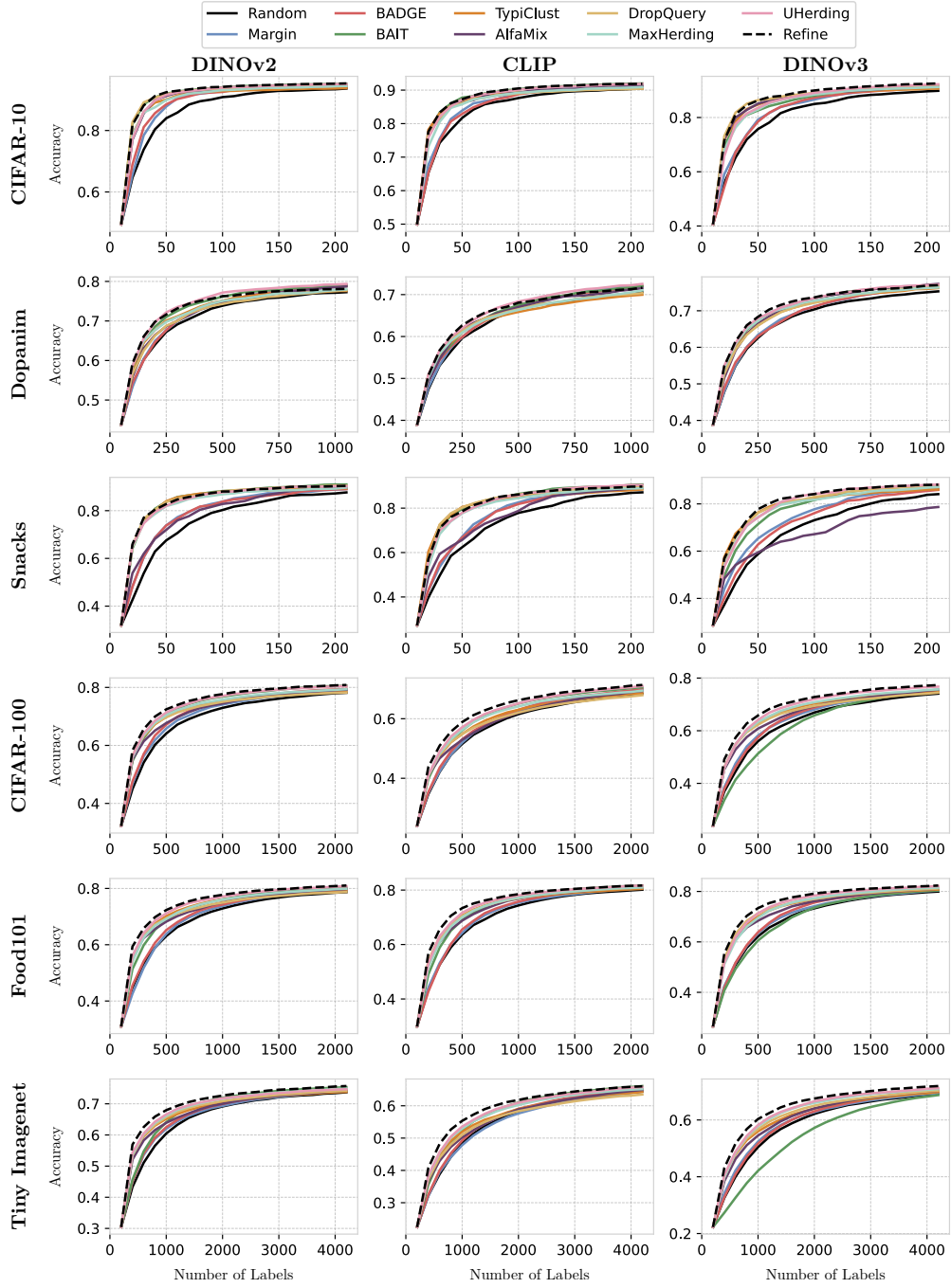


Figure 10. Complete set of absolute learning curves for individual AL strategies showing the accuracy for all datasets and backbone combinations.

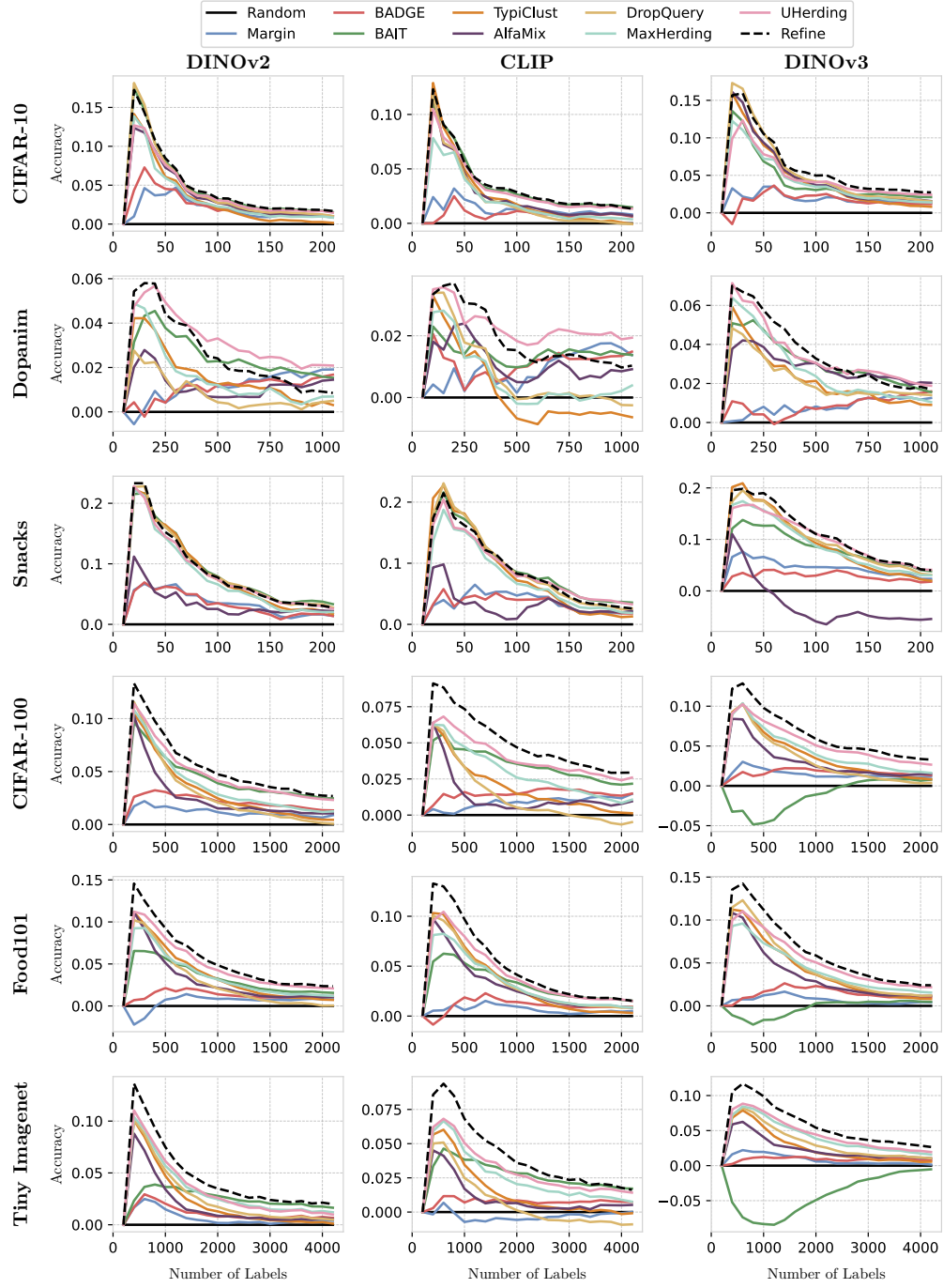


Figure 11. Complete set of relative learning curves for individual AL strategies showing the accuracy relative to random sampling from $\mathcal{U}_t$ for all datasets and backbone combinations.

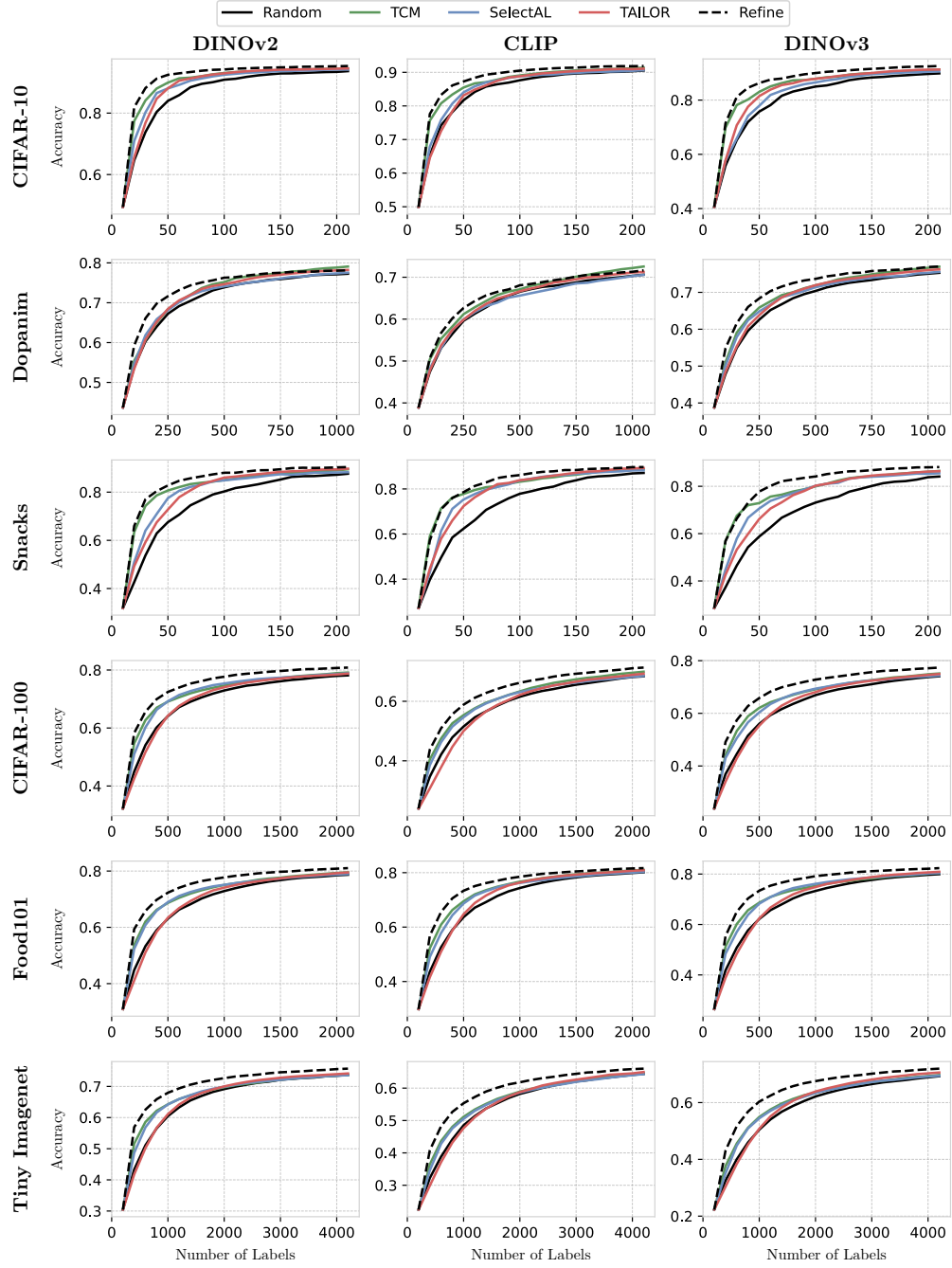


Figure 12. Complete set of absolute learning curves for ensemble AL methods showing the accuracy for all datasets and backbone combinations.

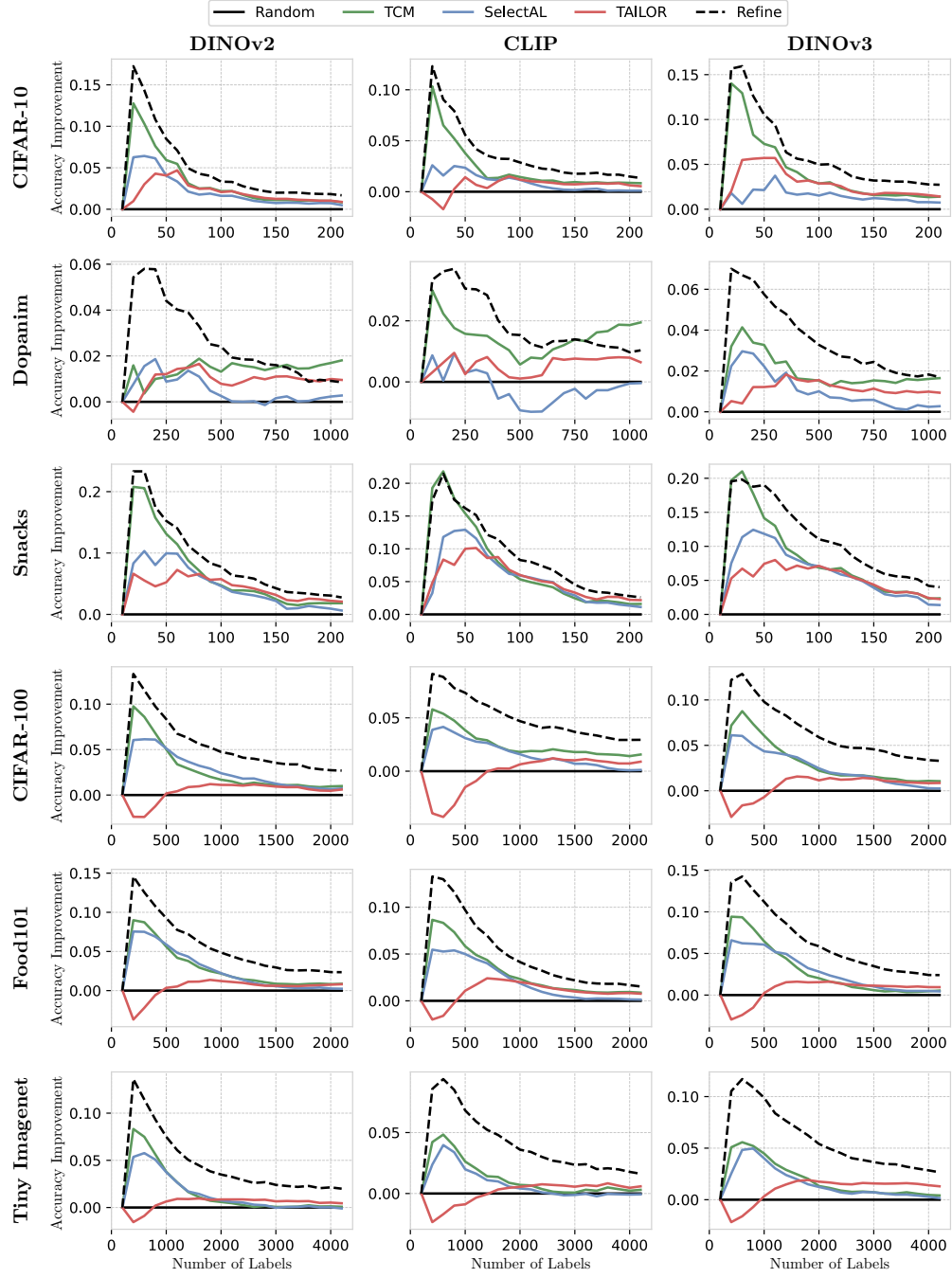


Figure 13. Complete set of relative learning curves for ensemble AL methods showing the accuracy relative to random sampling from $\mathcal{U}_t$ for all datasets and backbone combinations.