

Dynamics of Learning under User Choice: Overspecialization and Peer-Model Probing

Adhyyan Narang[†] Sarah Dean[‡] Lillian J. Ratliff[†] Maryam Fazel[†]

Electrical and Computer Engineering, University of Washington[†]
Computer Science, Cornell University[‡]

Correspondence: adhyyan@uw.edu

Abstract

In many economically relevant contexts where machine learning is deployed, multiple platforms obtain data from the same pool of users, each of whom selects the platform that best serves them. Prior work in this setting focuses exclusively on the “local” losses of learners on the distribution of data that they observe. We find that there exist instances where learners who use existing algorithms almost surely converge to models with arbitrarily poor global performance, even when models with low full-population loss exist. This happens through a feedback-induced mechanism, which we call the overspecialization trap: as learners optimize for users who already prefer them, they become less attractive to users outside this base, which further restricts the data they observe. Inspired by the recent use of knowledge distillation in modern ML, we propose an algorithm that allows learners to “probe” the predictions of peer models, enabling them to learn about users who do not select them. Our analysis characterizes when probing succeeds: this procedure converges almost surely to a stationary point with bounded full-population risk when probing sources are sufficiently informative, e.g., a known market leader or a majority of peers with good global performance. We verify our findings with semi-synthetic experiments on the MovieLens, Census, and Amazon Sentiment datasets.¹

1 Introduction

Traditional supervised learning theory typically assumes a single learner observing data drawn from a fixed distribution. However, this assumption is increasingly violated in modern machine learning markets, such as recommendation platforms and large language model (LLM) services. In these ecosystems, multiple learners operate on the same pool of users, and data is not assigned randomly. Instead, users choose which platform to engage with based on how well that platform serves their specific needs or preferences. Consequently, the data distribution observed by a learner is a function of the learner’s own performance and the choices available in the market. This setting is increasingly garnering interest in the machine learning community [9, 14, 16, 45, 47].

This coupling between model performance and user selection creates a feedback loop. As a learner optimizes for its current user base, it becomes increasingly specialized to that subpopulation. While this minimizes “local” loss on observed users, it often degrades performance on the unobserved population, a phenomenon we term *overspecialization*. Once a learner is overspecialized, it gets caught in an informational trap: it cannot learn to serve new users because it never observes them, and it never observes them because it cannot serve them. At a societal level, this dynamic fuels the formation of algorithmic echo chambers [4, 13, 26, 31], where platforms fragment the population rather than learning a robust, globally capable model.

Independently, another trend has become relevant in modern machine learning systems that has implications for the overspecialization problem: techniques such as knowledge distillation and training on synthetic data are becoming ubiquitous, particularly in the training of Large Language

¹Code for this paper is available at: <https://github.com/AdhyyanNarang/overspecialization-probing>.

Models [24, 52]. While these methods are typically employed to improve reasoning capabilities or computational efficiency (through compression of data), they introduce a structural change to the learning dynamic. Models are no longer limited to learning from organic user data, but can also "probe" other models to acquire synthetic labels. This enables learners to observe signals outside their siloed data distributions. In this work, we study whether the use of probing mechanisms in these machine learning markets can mitigate overspecialization.

Our Contributions. We model a market where users select learners based on a combination of inherent preferences and predictive loss. We analyze the resulting dynamics through a game-theoretic lens to understand the impact of peer probing. Our main findings are as follows:

1. **The Failure of Standard Learning:** We first analyze standard Multi-learner Streaming Gradient Descent (MSGD) in the absence of probing [47]. We prove that due to the user selection mechanism, MSGD can converge to "bad" stationary points. In these equilibria, learners become overspecialized, achieving low loss on their niche but arbitrarily poor performance on the global population.
2. **Convergence of Peer Probing:** We propose a new algorithm, **MSGD with Probing (MSGD-P)**, where learners mix gradient updates from organic users with updates from pseudo-labeled queries sent to peer models. We prove that this multi-agent dynamic converges to a stationary point of a modified potential function (Theorem 3). To our knowledge, we are the first to observe and study the multi-agent dynamics that arise from synthetic data training.
3. **Restoring Global Competence:** We show that probing effectively mitigates overspecialization. By characterizing the stationary points of MSGD-P, we derive bounds on the full-population loss under appropriate informational conditions (e.g., probing a known market leader vs. aggregating diverse peers).
4. **Empirical Validation:** We validate our findings on the MovieLens, US Census and Amazon Sentiment datasets. We observe that while standard learning leaves some models trapped with poor accuracy, introducing peer probing closes this performance gap.

2 Related Work

Our work sits at the intersection of three lines of research. We study a *multi-learner* setting where users select among competing platforms based on a combination of *inherent preferences and predictive quality*: a model richer than pure loss-minimization or uniform-random selection. Motivated by modern distillation practices, we analyze *peer-model probing* as a mechanism to mitigate overspecialization.

Performative Prediction. Our setting is an instance of performative prediction [20, 37, 40], where deploying a model influences the data a learner subsequently observes. Multi-learner extensions [17, 33, 38, 41, 50, 51, 58] provide general tools for endogenous distribution shift. We specialize to user-choice, where shift arises from selection rather than manipulation, enabling precise characterization of overspecialization and peer-model probing.

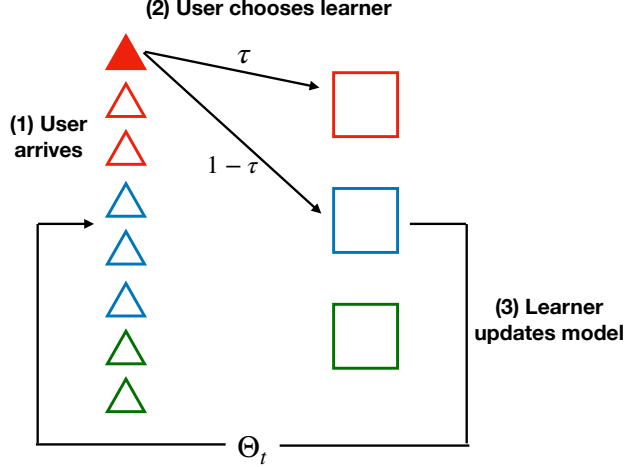


Figure 1: Illustration of our online multi-learner problem setting. The borders of users represent their highest ranked learner $\pi(z)$. For further details, see Section 3.

Learning under User Agency. Prior work on user agency focuses on single-learner settings where users opt out or selectively provide data [5, 10, 22, 23, 30, 42, 56]; these analyses do not capture the inter-learner feedback dynamics that arise in our multi-learner setting. In the multi-learner setting, Ginart et al. [16] characterize the *existence* of performance gaps due to competition via batch retraining; we prove that streaming dynamics *converge to* overspecialized equilibria (Theorem 2) and propose probing as a mitigation. Most closely related, Dean et al. [14] and Su and Dean [47] study gradient-based dynamics in choice-driven settings; we extend their framework to analyze full-population risk and introduce peer probing. See Appendix A for detailed comparisons with Ginart et al. [16], Kwon et al. [32], Bose et al. [9], and Su and Dean [47].

Knowledge Distillation. Our probing mechanism draws on knowledge distillation [24] and self-training [44, 55]; unlike Deep Mutual Learning [57] and codistillation [2], which assume shared data, our setting couples distillation with user-driven selection, yielding novel multi-agent dynamics. This also distinguishes our setting from domain generalization methods, which train a centralized model from labeled source domains, and selective-label/sample-selection work, which typically studies a single learner under an exogenous censoring rule; here each learner’s observed distribution is endogenously determined by user choices among competing models.

3 Problem Setting

We now formalize the machine learning market described above.

3.1 Users and Learners

Consider a market with m service providers (learners) serving a population of users distributed according to $\mathcal{P}$ over $\mathcal{Z} = \mathcal{X} \times \mathcal{Y}$, where $\mathcal{X} \subseteq \mathbb{R}^d$ denotes covariates and $\mathcal{Y}$ denotes labels ($\mathcal{Y} \subseteq \mathbb{R}$ for regression, $\mathcal{Y} \subseteq \{1, \dots, C\}$ for classification). We write $\mathcal{P}_X$ for the marginal distribution over covariates and $\mathcal{P}_{Y|X}(\cdot | x)$ for the conditional label distribution. When densities exist, we denote them by p_X and $p_{Y|X}$.

Each learner $i \in [m] := \{1, \dots, m\}$ maintains a model with parameters θ_i , and we write $\Theta = (\theta_1, \dots, \theta_m)$ for the joint parameter vector. The loss $\ell(x, y; \theta)$ measures the cost incurred by a model with parameters θ on a user with features x and label y . For regression, we consider linear predictors $h_\theta(x) = x^\top \theta$ with squared loss

$$\ell_{\text{SQ}}(x, y; \theta) = (y - x^\top \theta)^2.$$

For classification with C classes, we use cross-entropy loss

$$\ell_{\text{CE}}(x, y; \theta) = - \sum_{c=1}^C y_c \log h_\theta(x)_c,$$

where $h_\theta(x)_c = e^{x^\top \theta_c} / \sum_{j=1}^C e^{x^\top \theta_j}$ is the softmax output. Our framework can handle nonlinear relationships as well, as long as a suitable nonlinear feature transformation is known i.e if $h_\theta(x) = \phi(x)^\top \theta$.

Assumption 1. *The distribution $\mathcal{P}$ has continuous density $p_Z(z) = p_X(x)p_{Y|X}(y|x)$ with p_X supported on $\{x : \|x\| \leq R\}$ for some $R > 0$. For regression, $p_{Y|X}$ is supported on $[-Y_{\max}, Y_{\max}]$.*

Under this assumption, both loss functions satisfy standard regularity conditions (see Appendix B).

Lemma 1. *Under Assumption 1, for all $z \in \mathcal{Z}$, the loss $\ell(z, \cdot)$ is non-negative, convex, differentiable, locally Lipschitz, and β_ℓ -smooth.*

3.2 User Preferences and Platform Choice

A central feature of our framework is that users have *inherent preferences* over platforms that exist independently of current model quality. These preferences capture factors such as brand loyalty, familiarity, network effects, or historical habits. We encode them via a function $\pi : \mathcal{Z} \rightarrow [m]$, where $\pi(z)$ denotes the platform that user z intrinsically prefers. This function is exogenous and fixed throughout learning: it represents pre-existing affinities that platforms cannot directly influence through model updates.

The preference function π induces a partition of the user space. Let $S_i = \{z \in \mathcal{Z} : \pi(z) = i\}$ denote users who prefer platform i , with $\mathcal{P}_i = \mathcal{P}|_{S_i}$ the corresponding conditional distribution and $\alpha_i = \Pr_{z \sim \mathcal{P}}[\pi(z) = i]$ the fraction of users preferring platform i . Users do not only follow their inherent preferences; they also consider predictive quality. We model this tradeoff as follows:

Definition 1 (User Selection Rule). *Given models Θ , user z selects platform*

$$M(z; \Theta) = \begin{cases} \pi(z) & \text{with probability } \tau, \\ \arg \min_{i \in [m]} \ell(z; \theta_i) & \text{with probability } 1 - \tau. \end{cases}$$

The parameter $\tau \in [0, 1]$ governs how strongly inherent preferences influence user behavior. When $\tau = 1$, users ignore model quality entirely and follow their intrinsic preferences. When $\tau = 0$, users select the platform minimizing their loss. For intermediate values, users often default to familiar platforms but sometimes shop based on quality. This model generalizes the setting of Su and Dean [47], who assume users either minimize loss or choose uniformly at random. We make this choice because user preferences in practice are neither purely quality-driven nor purely random: the $\pi(z)$ function captures persistent individual affinities that create systematic heterogeneity in platform selection.

Model quality also induces a partition of users. Let $Z_i(\Theta) = \{z : i = \arg \min_{j \in [m]} \ell(z; \theta_j)\}$ denote users for whom platform i achieves minimal loss, with $\mathcal{D}_i(\Theta) = \mathcal{P}|_{Z_i(\Theta)}$ the conditional distribution and $a_i(\Theta) = \Pr_{z \sim \mathcal{P}}[z \in Z_i(\Theta)]$ the population fraction. Under Definition 1, learner i observes users from the mixture

$$\mathcal{O}_i(\Theta) = \tau \alpha_i \mathcal{P}_i + (1 - \tau) a_i(\Theta) \mathcal{D}_i(\Theta),$$

where $\mathcal{O}_i(\Theta)$ is a sub-probability measure with total mass $w_i(\Theta) = \tau \alpha_i + (1 - \tau) a_i(\Theta)$.

3.3 Dynamics and Learning Objective

We consider an online setting where learners interact with users over T timesteps, illustrated in Figure 1. At each step t :

1. A user $z^t \sim \mathcal{P}$ arrives and selects platform $M(z^t; \Theta^t)$.
2. The selected learner observes z^t , incurs loss $\ell(z^t; \theta_{M(z^t; \Theta^t)}^t)$, and updates its parameters.

We denote an instance of this market as $\mathcal{G} = (\mathcal{P}, \ell, \pi, \tau)$.

Learning Objective. Each learner’s goal is to minimize the *full-population risk*

$$\mathcal{R}(\theta) = \mathbb{E}_{z \sim \mathcal{P}}[\ell(z; \theta)].$$

We write $\theta^* = \arg \min_{\theta} \mathcal{R}(\theta)$ for the population-optimal model and $\epsilon = \mathcal{R}(\theta^*)$ for the Bayes risk. This objective differs from prior work [14, 47], which focuses on the “local” loss over each learner’s observed distribution $\mathcal{O}_i(\Theta)$. We focus on full-population risk because the signature of overspecialization is precisely the gap between local and global performance: a learner may achieve low loss on $\mathcal{O}_i(\Theta)$ while performing poorly on users outside its observed base. Understanding when this gap emerges and how to prevent it is the central question of this paper.

4 The Failure of Standard Learning Dynamics

We now analyze Multi-learner Streaming Gradient Descent (MSGD), the standard algorithm for this setting introduced by Su and Dean [47]. First, we reproduce the result from [47] that MSGD converges to stationary points (Theorem 1); however, we use a different proof technique that we believe is more insightful, which we describe below. Despite convergence, we show that these equilibria can exhibit severe overspecialization: learners achieve low loss on observed users while performing arbitrarily poorly on the full population (Theorem 2).

4.1 Algorithm and Assumptions

Algorithm 1 presents MSGD, studied in prior work by Su and Dean [47]. When a user arrives and selects learner i , that learner performs a stochastic gradient update on the observed loss, and other learners remain unchanged.

In order to study the convergence behavior of the algorithm, we make the following standard assumptions on learning rates and loss geometry, which are the same as in Su and Dean [47].

Assumption 2. *The learning rates satisfy $\sum_{t=1}^{\infty} \eta^t = \infty$ and $\sum_{t=1}^{\infty} (\eta^t)^2 < \infty$.*

Assumption 3. *For any $\theta \neq \theta'$, there exists $d_0 > 0$ such that for all $d < d_0$, the set $\{z : |\ell(z; \theta) - \ell(z; \theta')| < d\}$ has Lebesgue measure at most d .*

Algorithm 1 Multi-learner Streaming Gradient Descent (MSGD) [47]

Require: Loss function $\ell(\cdot, \cdot) \geq 0$; initial models $\Theta^0 = (\theta_1^0, \dots, \theta_m^0)$; learning rate $\{\eta^t\}_{t \geq 1}$

- 1: **for** $t = 0, 1, 2, \dots, T$ **do**
- 2: Sample user $z^t \sim \mathcal{P}$
- 3: User selects learner $i = M(z^t; \Theta^t)$
- 4: $\theta_i^{t+1} \leftarrow \theta_i^t - \eta^t \nabla_{\theta} \ell(z^t; \theta_i^t)$
- 5: **end for**
- 6: **return** Θ^T

Assumption 3, from Su and Dean [47], rules out pathological cases where a large mass of users is nearly indifferent between two learners. Intuitively, it ensures that small parameter perturbations move only a small mass of users across loss-induced decision boundaries, so the partition-dependent potential f can be differentiated without boundary terms. Assumption 4 below is the standard boundedness condition used in stochastic-approximation analyses.

4.2 Convergence to Stationary Points

In MSGD, each learner i optimizes its expected loss over observed users. Define the potential function $f(\Theta)$ as the sum of the expected losses of all learners on the distributions that they observe.

$$f(\Theta) = \sum_{i=1}^m \mathbb{E}_{\mathcal{O}_i(\Theta)}[\ell(z; \theta_i)], \quad (1)$$

where we recall that $\mathcal{O}_i(\Theta) = \tau \alpha_i \mathcal{P}_i + (1 - \tau) a_i(\Theta) \mathcal{D}_i(\Theta)$, and $w_i(\Theta) = \tau \alpha_i + (1 - \tau) a_i(\Theta)$. Note that this is the sum of the “local” losses of each of the learners on their observed distribution, unlike the global risk $\mathcal{R}(\cdot)$ that was introduced above.

In order to show convergence, we make the same boundedness assumptions as in Su and Dean [47].

Assumption 4. *The parameter sequence $\{\Theta^t\}_{t \geq 0}$ is almost surely bounded: $\sup_{t \geq 0} \|\Theta^t\| < \infty$. Moreover, the set $\{\Theta : \nabla f(\Theta) = 0\}$ is compact.*

Our approach uses stochastic approximation [8] and differs from that of Su and Dean [47]. The key insight is that f serves as a Lyapunov function: despite each learner optimizing over a different, endogenously-determined user distribution, the aggregate of local losses forms a coherent potential. This is surprising because multi-agent gradient dynamics often cycle or diverge [36]; the Lyapunov structure explains why convergence occurs here. Concretely, we show that the MSGD iterates track the ODE $\dot{\Theta} = -\nabla f(\Theta)$. The following lemma, proved using standard stochastic approximation arguments, establishes this connection.

Lemma 2. *Let Assumptions 1–4 hold. Then the iterates $\{\Theta^t\}$ of Algorithm 1 converge to a compact connected internally chain transitive invariant set of the ODE $\dot{\Theta} = -\nabla f(\Theta)$.*

See Appendix B.1 for definitions of the dynamical-systems terms used in Lemma 2. The lemma shows that the discrete stochastic dynamics behave, in the limit, like the continuous gradient flow on f . Since f decreases along trajectories of this ODE i.e. $\frac{d}{dt} f(\Theta(t)) = -\|\nabla f\|^2 \leq 0$, the only invariant sets are stationary points. This yields our main convergence result.

Theorem 1. *Let Assumptions 1–4 hold. Then the iterates $\{\Theta^t\}$ of Algorithm 1 converge to the set of stationary points $\{\Theta : \nabla f(\Theta) = 0\}$ almost surely.*

Hence, the stochastic approximation framework reveals that MSGD implicitly minimizes a single global objective f , providing a clean explanation for why convergence occurs despite the multi-agent structure.

The same potential argument extends to mini-batch updates, matching the implementation used in the experiments. If a batch $\mathcal{B}^t$ is sampled i.i.d. and learner i updates using the average gradient over the selected subset $S_i^t = \{z \in \mathcal{B}^t : M(z; \Theta^t) = i\}$, then conditional on $S_i^t \neq \emptyset$ this average is an unbiased gradient over $\mathcal{O}_i(\Theta^t)/w_i(\Theta^t)$. Since $w_i(\Theta) \geq \tau\alpha_i > 0$ for $\tau > 0$, the batch dynamics track the same potential with a bounded state-dependent rescaling of the step size.

4.3 The Overspecialization Trap

While Theorem 1 guarantees convergence, the stationary points may be highly undesirable. A learner cannot improve on users it never observes, and it never observes users it cannot serve well. This feedback loop is the overspecialization trap.

The following theorem shows that this trap can be severe: MSGD can converge to equilibria where some learners have arbitrarily poor global performance, even when models with low full-population loss exist.

Theorem 2. *Let Assumptions 1–4 hold. For any $\tau \geq \frac{1}{2}$ and any choice of ϵ, Γ with $0 < \epsilon < \Gamma$, there exists an instance $\mathcal{G}$ such that:*

1. *There exists θ^* with $\mathcal{R}(\theta^*) \leq \epsilon$.*
2. *The MSGD iterates converge to a unique stationary point $\bar{\Theta}$ where $\mathcal{R}(\bar{\theta}_i) \geq \Gamma$ for some learner $i \in [m]$.*

The key mechanism is as follows. When $\tau \geq \frac{1}{2}$, inherent preferences dominate at equilibrium: the loss-induced partition collapses to the ranking partition, so that $Z_i(\Theta) = S_i$ for all i . Each learner optimizes exclusively for users who intrinsically prefer it, arriving at

$$\bar{\theta}_i = \arg \min_{\theta} \mathbb{E}_{z \sim \mathcal{P}_i} [\ell(z; \theta)].$$

This specialization occurs regardless of whether a better global model exists. In the constructed instance (detailed in Appendix C), learner 1 achieves *zero* loss on its observed population $\mathcal{P}_1$ while its full-population loss exceeds Γ . The learner has perfectly fit its niche while becoming arbitrarily poor globally. This dynamic formalizes the echo chamber phenomenon that platforms become increasingly specialized to their existing audience, unable to learn models that serve the broader population.

The condition $\tau \geq \frac{1}{2}$ delineates the regime where inherent preferences dominate user behavior. A key contribution of the analysis is showing that in this regime the equilibrium is unique and obtainable in closed form, which enables us to exactly characterize the risk. However, when $\tau < \frac{1}{2}$ and quality-based selection dominates, there may be multiple equilibria. Now, the limit point becomes initialization-dependent, which presents a technical obstacle to characterizing the limiting risk in closed form; however, experiments in Section 6 demonstrate similar phenomena across different values of τ .

5 Mitigating Overspecialization through Peer Probing

The standard MSGD dynamics from [47] fail to converge to models that perform well on the full population because they are “blocked” from observing users outside the ones who choose them, and are limited to learning from a restricted portion of the feature space.

However, in many practical settings, learners can *probe* the predictions of other learners. For example, a platform could create an account on a competitor’s service to observe their recommendations. This concept has gained particular prominence in the context of Large Language Models (LLMs) through *knowledge distillation* [24], where one model learns from another model’s outputs. Recent work has extensively explored this area [3, 34, 48, 52–54]. The release of the DeepSeek LLM [52] has brought knowledge distillation to the forefront of both policy discussions and mainstream media attention [52]. Despite this growing practical importance, the theoretical implications of these multi-agent learning interactions—where models train on each other’s outputs—remain largely unexplored. This section asks the question: under what circumstances can probing the predictions of others help overcome the overspecialization trap?

5.1 Algorithm

We propose MSGD with Probing (MSGD-P), shown in Algorithm 2. The algorithm has two phases. In the *offline phase*, each probing learner $j \in U$ collects a dataset $\mathfrak{D}_j$ of pseudo-labeled examples by querying peer models on sampled covariates. In the *online phase*, learners interleave standard MSGD updates (on organic users) with gradient steps on their probing datasets.

More concretely, for each probing learner $i \in U$, we sample covariates $(\tilde{x}_i^1, \dots, \tilde{x}_i^n) \sim \mathcal{P}_X^n$. Given a query covariate x , the learner selects a subset of peers $T_i(x) \subseteq [m]$ to consult, and forms a pseudo-label via median aggregation

$$y_{\text{agg},i}(x, \Theta) = \text{median}\{h_{\theta_j}(x) : j \in T_i(x)\}.$$

We then define $\tilde{y}_i^q := y_{\text{agg},i}(\tilde{x}_i^q, \Theta^0)$ and the pseudo-labeled examples $\tilde{z}_i^q := (\tilde{x}_i^q, \tilde{y}_i^q)$, and collect the probing dataset $\mathfrak{D}_i := \{\tilde{z}_i^q\}_{q=1}^n$. The choice of $T_i(x)$ determines which peers are consulted; we discuss this in Section 5.3.

For a probing learner $i \in U$, the update can be interpreted as a stochastic gradient step on the following instantaneous loss:

$$L_i^t(\theta_i) = \tau \alpha_i \mathbb{E}_{z \sim \mathcal{P}_i} [\ell(z; \theta_i)] + (1 - \tau) \alpha_i (\Theta^t) \mathbb{E}_{z \sim \mathcal{D}_i(\Theta^t)} [\ell(z; \theta_i)] + \frac{p}{n} \sum_{q=1}^n \ell(\tilde{z}_i^q; \theta_i) + \frac{\lambda p}{2} \|\theta_i\|^2. \quad (2)$$

The first two terms capture organic learning from users who select learner i (via inherent preference $\mathcal{P}_i$ or quality-based choice $\mathcal{D}_i$), while the third term captures learning from probing data. The parameter $p > 0$ controls the relative weight of probing gradients: larger p emphasizes pseudo-labels, smaller p prioritizes organic data.

We assume that probing learners can sample covariates from the full distribution $\mathcal{P}_X$, but do not have access to true labels. This asymmetry is natural in practice: covariates are often publicly available or easy to generate, while labels require costly human annotation or reveal private user behavior. For example, a streaming service knows the metadata of all movies (genre, cast, runtime) and general user demographics, but not which movies a non-subscriber would rate highly. In the LLM setting, covariates are simply text prompts, which can be generated synthetically or scraped from public sources (e.g., Reddit, StackOverflow), whereas ground-truth responses require expensive human evaluation.

We focus on *offline probing*, where pseudo-labels are collected once at initialization from a fixed snapshot of peer models. This design choice mirrors practical knowledge distillation workflows: in settings like DeepSeek distilling from Claude or GPT-4, teachers are queried to create a fixed dataset, and student training then proceeds independently [52]. Offline probing also ensures reproducibility by capturing a specific model version’s behavior, avoiding inconsistencies from querying peers at different stages of adaptation. We discuss extensions to online probing in Section 7.

Algorithm 2 Multi-learner Streaming Gradient Descent with Probing (MSGD-P)

Require: loss function $\ell(\cdot, \cdot) \geq 0$; Initial models $\Theta^0 = (\theta_1^0, \dots, \theta_m^0)$; Learning rate $\{\eta^t\}_{t=1}^{T+1}$, probing weight $p > 0$, set of probing learners $U \subseteq [m]$, regularization weight $\lambda \geq 0$

```
1: // Offline probing data collection
2: for  $j \in U$  : do
3:   Sample covariates  $(\tilde{x}_j^1, \dots, \tilde{x}_j^n) \sim \mathcal{P}_X^n$ .
4:   Collect pseudo-labels and store dataset  $\mathfrak{D}_j = \{(\tilde{x}_j^1, y_{\text{agg},i}(\tilde{x}_j^1, \Theta^0)) \dots (\tilde{x}_j^n, y_{\text{agg},i}(\tilde{x}_j^n, \Theta^0))\}$ 
5: end for
6: // Online updates
7: for  $t = 0, 1, 2, \dots, T$  do
8:   Sample data point  $z^t \sim \mathcal{P}$ 
9:   User selects model  $i = M(z^t; \Theta^t)$ 
10:   $\theta_i^{t+1} \leftarrow \theta_i^t - \eta^t \nabla \ell(z^t, \theta_i^t)$ 
11:  for  $j \in U$  : do
12:    Sample  $\tilde{z}_j^t$  uniformly from  $\mathfrak{D}_j$ 
13:     $\theta_j^{t+1} \leftarrow \theta_j^t - \eta^t p \left( \nabla \ell(\tilde{z}_j^t, \theta_j^t) + \lambda \theta_j^t \right)$ 
14:  end for
15: end for
16: return  $\Theta^T$ 
```

5.2 Convergence

With probing, each learner $i \in U$ now optimizes a blend of two objectives: the loss on observed users (as in standard MSGD) and the loss on probing data. This leads to a modified potential function:

$$\tilde{f}(\Theta) = f(\Theta) + p \sum_{i \in U} \left(\frac{1}{n} \sum_{q=1}^n \ell(\tilde{z}_i^q, \theta_i) + \frac{\lambda}{2} \|\theta_i\|^2 \right), \quad (3)$$

where the second term captures the probing loss (with regularization) weighted by p .

The same stochastic approximation analysis from Section 4 extends to this setting: $\tilde{f}$ serves as a Lyapunov function for the modified dynamics, yielding the following convergence guarantee.

Theorem 3. *Let Assumptions 1-3, Assumption 4 (as applied to $\tilde{f}$), and Assumption 5 hold. Then, the iterates $\{\Theta^t\}$ of Algorithm 2 converge to the set of stationary points $\{\Theta : \nabla \tilde{f}(\Theta) = 0\}$ almost surely.*

Crucially, the stationary points of $\tilde{f}$ differ from those of f : probing changes *where* learners converge, not *whether* they converge. The probing term pulls learners toward models that perform well on the probed distribution, potentially escaping the overspecialization trap. However, this benefit depends on the quality of the pseudo-labels—a question we address next.

5.3 When Does Probing Help?

Note that the convergence guarantee above holds for any choice of $T_i(x)$. However, in order for the probing data to be *helpful*, the pseudo-labels must be a good proxy for the ground-truth labels.

Scenario	What i knows	Peer requirement	$T_i(x)$	B
Majority-good	Nothing	$> 50\%$ in $B_r(\theta^*)$	$[m]$	$R^2 r^2 + 2\epsilon$
Market-leader	Identity of j^*	$\mathcal{R}(\theta_{j^*}^0) \leq \xi$	$\{j^*\}$	ξ
Partial knowledge	Subset G	$> 50\%$ of G in $B_r(\theta^*)$	G	$R^2 r^2 + 2\epsilon$
Preference-aware	$\pi(x)$	Nothing	$\{\pi(x)\}$	ϵ

Table 1: Probing scenarios with corresponding rules and accuracy bounds. The preference-aware scenario is notable: it requires no assumption on peer quality, only knowledge of user preferences.

Assumption 5 (Accurate Probing). *We say that the “accurate probing” condition holds for probing learner $i \in [m]$ if there exists $B \geq 0$ such that*

$$\mathbb{E}_{(x,y) \sim \mathcal{P}} \left[(y_{agg,i}(x, \Theta_{-i}) - y)^2 \right] \leq B.$$

This assumption is stated for a single probing learner i . Different learners may satisfy it via different scenarios (or not at all); performance guarantees in Section 5.4 apply to any learner for whom the assumption holds.

Below, we identify scenarios under which Assumption 5 holds. These scenarios differ along two axes: *what learner i must know* about the market, and *what must be true* about peer models at initialization. Table 1 summarizes these scenarios; the second and third columns display the knowledge-vs-peer-requirement tradeoff.

Definition 2 (Globally good peers). *We say learner i can achieve accurate probing via globally good peers if the following instance properties and knowledge conditions hold:*

- (i) **Majority-good.** *More than half of the learners satisfy $\theta_j^0 \in B_r(\theta^*)$ for a given proximity parameter $r > 0$.*
- (ii) **Market-leader.** *There exists a single learner $j^* \in [m]$ such that*

$$\mathbb{E}_{x \sim \mathcal{P}} [\ell(x, \theta_{j^*}^0)] \leq \xi,$$

and the identity of j^ is known to learner i .*

- (iii) **Partial knowledge.** *There exists a subset $G \subseteq [m] \setminus \{i\}$ such that more than half of the learners in G satisfy $\theta_j^0 \in B_r(\theta^*)$ for a given $r > 0$, and learner i has knowledge of the subset G .*

Above, the parameter $r > 0$ in the majority-good and partial-knowledge scenarios controls how close peers must be to θ^* ; smaller r yields tighter bounds. When no globally good learners exist or can be identified, probing may still be effective if the learner has access to ranking information.

Definition 3 (Preference-aware probing). *Suppose all learners initialize at the parameters $\bar{\Theta} = (\bar{\theta}_1, \dots, \bar{\theta}_m)$. If learner i has knowledge of the inherent preference function $\pi(z)$, which identifies each user’s preferred platform, we call this the preference-aware scenario.*

These definitions capture different approaches to probing: *Definition 2* covers settings where learners can identify or rely on peers with strong global performance, while *Definition 3* addresses settings where learners must instead leverage knowledge of user preferences. The scenarios exhibit a fundamental tradeoff: when peer models are favorable (e.g., many are globally good), learner i needs

little knowledge to achieve accurate probing; conversely, with stronger knowledge (e.g., knowing $\pi(x)$), the learner can probe in a more targeted way, and probing succeeds even when no peer is globally competent.

In realistic markets, learners often naturally gain access to information enabling one of these approaches. For Scenarios ((i)) or ((iii)), platforms may observe broad industry benchmarks [7, 12, 35, 39, 43]. For Definition 3, it is natural to maintain knowledge of user preference patterns [1, 18, 25, 27]. For Scenario (ii), there are examples in the LLM literature of this explicitly happening: for instance, the Alpaca model [49] was admittedly trained on data generated by text-davinci-003 (GPT-3.5), the Vicuna model [11] was trained on data generated by GPT-4, and the survey paper [19] documents the prevalence of this practice.

The probing rule $T_i(x)$ (fourth column of Table 1) in each scenario is chosen so that median aggregation is robust: probing all peers when the majority are good, targeting the known leader when one exists, or routing to the locally-expert peer in the preference-aware case. The preference-aware scenario is particularly notable: it requires *no assumption* on peer quality, only knowledge of user preferences $\pi(x)$, enabling learner i to aggregate specialized knowledge into global competence even when every peer suffers from overspecialization.

The next lemma shows that in each scenario, the pseudo-labels obtained via the corresponding $T_i(x)$ are uniformly bounded in mean-squared error, ensuring that Assumption 5 holds.

Lemma 3. *For each scenario in Definitions 2–3, the probing rule $T_i(x)$ in Table 1 satisfies Assumption 5 with the stated accuracy parameter B .*

5.4 Performance Guarantees

We now characterize the full-population risk of MSGD-P stationary points. The bound decomposes into four interpretable components: an irreducible Bayes error term, a probing bias term capturing pseudo-label inaccuracy, and two regularization-dependent terms reflecting a bias-variance tradeoff.

Theorem 4. *Let Assumptions 1–5 hold. For any $\kappa \in (0, 1)$, with probability at least $1 - \kappa$, we have for probing learner $i \in [m]$, for the squared loss:*

$$\mathcal{R}(\tilde{\theta}_i) \leq O\left(\left(\frac{p+1}{p}\right)\epsilon + B + \lambda\|\theta^*\|^2 + \frac{(p+1)C_{gen}}{p\lambda}\sqrt{\frac{\log(1/\kappa)}{n}}\right).$$

where $C_{gen}(R, Y_{\max}, \|\theta^*\|, \max_{j \neq i} \|\theta_j^0\|)$ is stated explicitly in the Appendix.

The four terms admit natural interpretations: (i) $\frac{p+1}{p}\epsilon$ is the irreducible Bayes error, scaled by the ratio of total to probing gradient weight; (ii) B is the probing bias from pseudo-label inaccuracy (see Table 1); (iii) $\lambda\|\theta^*\|^2$ is the regularization bias; and (iv) the final term captures finite-sample generalization error from n probing queries.

The probing weight p controls the balance between the organic distribution (users who choose learner i) and the probing distribution (the full population). A learner who prioritizes performance on their existing user base may prefer smaller p , while a learner seeking global competence benefits from larger p .

The regularization parameter λ exhibits a classical bias-variance tradeoff. Larger λ increases the regularization bias ($\lambda\|\theta^*\|^2$) but improves generalization by keeping parameter norms bounded, reducing the $O(1/\sqrt{\lambda n})$ term. Conversely, smaller λ reduces bias but worsens the generalization bound.

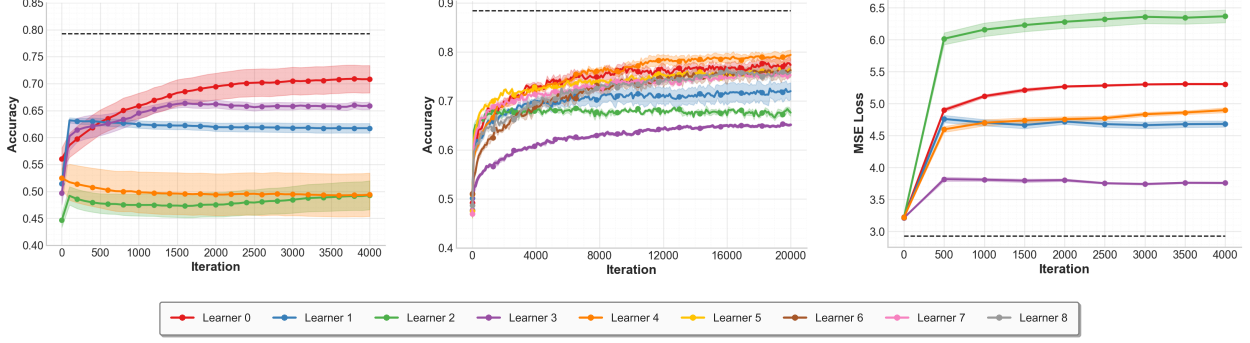


Figure 2: **MSGD full-population performance with random initialization (Preference-aware scenario)**. Left: Census test accuracy. Mid: Amazon sentiment test accuracy Right: MovieLens test loss. The dashed black line represents the performance of a baseline θ^* trained on the full dataset. In all cases, $\tau = 0.3$; other dataset-specific hyperparameters are given in Table 3.

Corollary 1. *Let Assumptions 1–5 hold. Fix $\lambda = \epsilon/\|\theta^*\|^2$ and any $\kappa \in (0, 1)$. Define $M_0 := \max_{j \neq i} \|\theta_j^0\|$. If there are sufficiently many probing samples $n \geq \underline{n}(p, \kappa, \epsilon, R, Y_{\max}, M_0, \|\theta^*\|)$, then with probability at least $1 - \kappa$, every stationary point $\tilde{\Theta}$ of MSGD-P satisfies, for each probing learner i ,*

$$\mathcal{R}(\tilde{\theta}_i) \leq O\left(\left(\frac{p+1}{p}\right)\epsilon + B\right).$$

See Appendix E for the explicit sample complexity $\underline{n}$ and proof. We note that this sample complexity bound is not tight in n : we show in Figure 5 that strong empirical recovery can occur with very small probing datasets.

Remark 1 (Cross-entropy loss). *An analogous performance guarantee holds for cross-entropy loss; see Assumption 6 and Appendix F for the bound, and Table 2 for the corresponding accuracy parameters.*

Hence, probing breaks the information barrier created by user-choice dynamics: a learner who is overspecialized cannot observe users outside its niche and thus cannot learn to serve them. While Theorem 2 shows that the loss of a learner may be arbitrarily worse than ϵ , for sufficiently large n , the bound above presents a ceiling on the risk of any probing learner.

6 Numerical Experiments

We evaluate our approach on three real-world datasets: MovieLens-10M, the ACS Employment dataset from the US Census (Alabama, 2018), and the Amazon Reviews 2023 corpus.

6.1 Experimental Setup

We evaluate on three datasets; in each case, every data point corresponds to an individual user.

Movie Recommendation with Squared Loss. We use the MovieLens-10M dataset [21], which contains 10 million movie ratings from 70k users across 10k movies—a natural testbed for multi-learner competition in recommendation. Following Bose et al. [9] and Su and Dean [47], we extract $d = 16$ dimensional user embeddings via matrix factorization and retain ratings for the top 200

most-rated movies, yielding a population of 69,474 users. Each user’s data consists of $z = (x, r)$ where $x \in \mathbb{R}^d$ is the embedding and r contains their ratings. Let Ω_x denote the set of movies rated by user x with $|\Omega_x|$ movies. Each learner fits a linear model $\theta \in \mathbb{R}^{d \times 200}$ using squared loss:

$$\ell(z; \theta) = \frac{1}{|\Omega_x|} \sum_{i \in \Omega_x} (\theta_i^\top x - r_i)^2.$$

Census Data with Logistic Loss. We use the ACSEmployment task from `folktables` [15], where the goal is to predict employment status from demographic features. The population consists of 38,221 individuals from the 2018 Alabama census (ages 16–90), with $d = 16$ features describing age, education, marital status, etc. Each user’s data is $z = (x, y)$ where $x \in \mathbb{R}^d$ (standardized to zero mean, unit variance) and $y \in \{0, 1\}$. Each learner uses logistic regression:

$$\ell(z; \theta) = -y \log(\sigma(\theta^\top x)) - (1 - y) \log(1 - \sigma(\theta^\top x)),$$

where σ is the sigmoid function. The model predicts $\hat{y} = \mathbf{1}[\theta^\top x > 0]$.

Amazon Reviews 2023 with Logistic Loss. We use the Amazon Reviews 2023 corpus and sample up to 30,000 reviews from nine product categories. Each review is represented by a $d = 384$ sentence embedding using `all-MiniLM-L6-v2`, and each learner performs binary sentiment classification with logistic regression. As in the Census setting, labels are binary and predictions are thresholded logistic outputs. Full preprocessing details are provided in Appendix G.

User Preferences. For Census and MovieLens ($m = 5$), we induce inherent preferences $\pi(z)$ via K-means clustering on user features, assigning each user to a preferred platform based on cluster membership. For Amazon ($m = 9$), preferences are determined by product category, with all reviews from a given category preferring the same learner.

Evaluation. For each dataset, the train/test split is fixed once and shared across all learners; there is no validation split, since no per-run hyperparameter tuning is performed. Census and Amazon report standard binary classification accuracy on the held-out test set, while MovieLens reports masked test MSE.

6.2 Experimental Results

We present results for the preference-aware scenario from Definition 3. The results are qualitatively similar in the other scenarios from Definition 2 as well, and are presented in Section G.

Expt 1: MSGD converges to equilibria with poor global performance. Our first set of experiments validates Theorem 2. Figure 2 shows the full-population performance trajectories of individual learners on all three datasets without probing ($p = 0$) over $T = 4000$ rounds, when initialized randomly. In all cases, the results reveal large overspecialization gaps relative to the dashed black baseline in Figure 2.

Expt 2: Peer Model Probing Mitigates Overspecialization. Figure 3 shows learner final performance at timestep T as a function of probing weight p using the offline probing budgets in Table 3; here, one learner (indicated by triangle markers) uses preference-aware probing while other learners use standard MSGD.

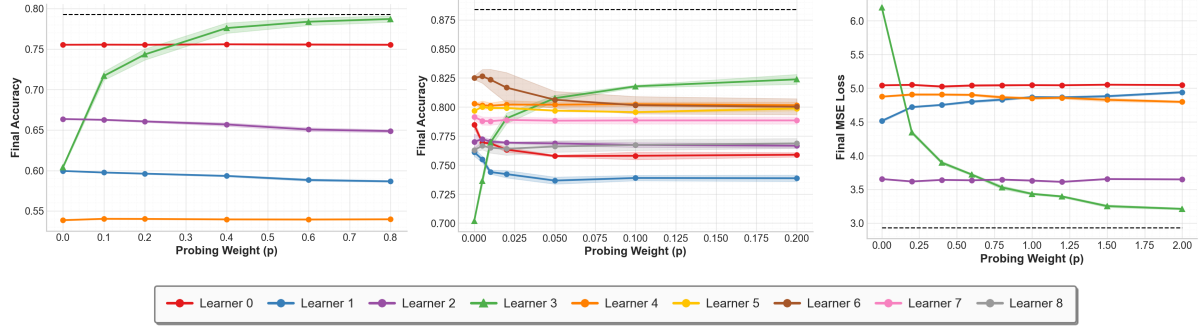


Figure 3: **Effect of probing on full-population performance (Preference-aware scenario).** Left: Census final accuracy vs probing weight p . Mid: Amazon sentiment final accuracy vs probing weight p . Right: MovieLens final loss vs p . In all cases, triangle markers indicate the probing learner. Here $\tau = 0.7$; other dataset-specific hyperparameters are given in Table 3.

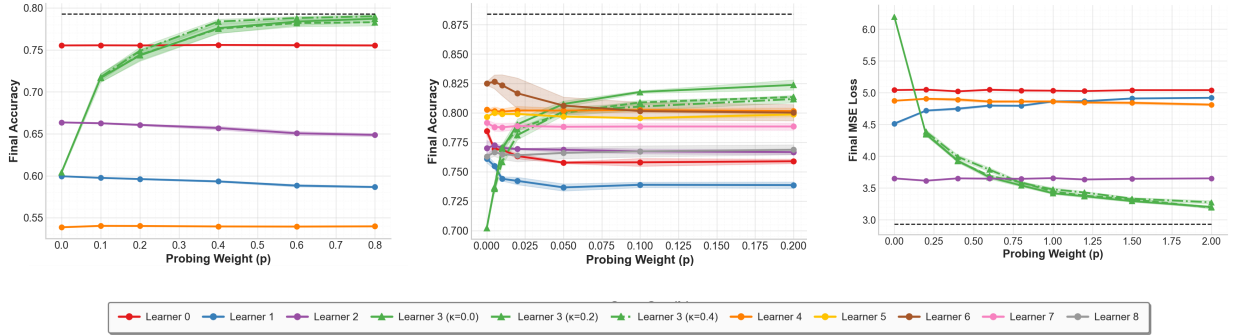


Figure 4: **Effect of noisy probing-source selection.** Preference-aware probing routes each probe query to the preferred peer $\pi(x)$ with probability $1 - \kappa$ and to a random peer with probability κ . Left: Census final test accuracy vs. probing weight p . Middle: Amazon sentiment final test accuracy vs. p . Right: MovieLens final test MSE vs. p . The green learner is the probing learner; its multiple curves show different κ values.

On Census (left), the probing learner’s accuracy improves from approximately 60% at $p = 0$ to about 78% at $p = 0.8$, shrinking its baseline gap from roughly 18 percentage points to about 1 percentage point. The improvement is monotonic in p : even modest probing ($p = 0.2$) yields noticeable gains. On MovieLens (right), the effect is equally pronounced: the probing learner’s MSE loss decreases from approximately 6.2 to 3.5.

Expt 3: Probing is robust to noisy source selection. Figure 4 tests a noisy version of preference-aware probing: for each probe query, the learner routes to $\pi(x)$ with probability $1 - \kappa$ and to a random peer with probability κ . Across datasets, probing remains beneficial under moderate routing noise, showing that the method does not require perfect knowledge of user preferences.

Expt 4: How much probing data is needed? Figure 5 studies sample efficiency by sweeping the probing dataset size n on Census, averaged over 10 random seeds. We observe substantial gains even with very small probing sets: for the probing learner with $p \in \{0.5, 1.0\}$, final accuracy rises from about 0.68 at $n = 5$ to about 0.78 by $n = 50$, and then saturates near 0.79 at $n = 100$; this is a tiny fraction of the full dataset size of 38,221 examples. As n increases, mean performance improves

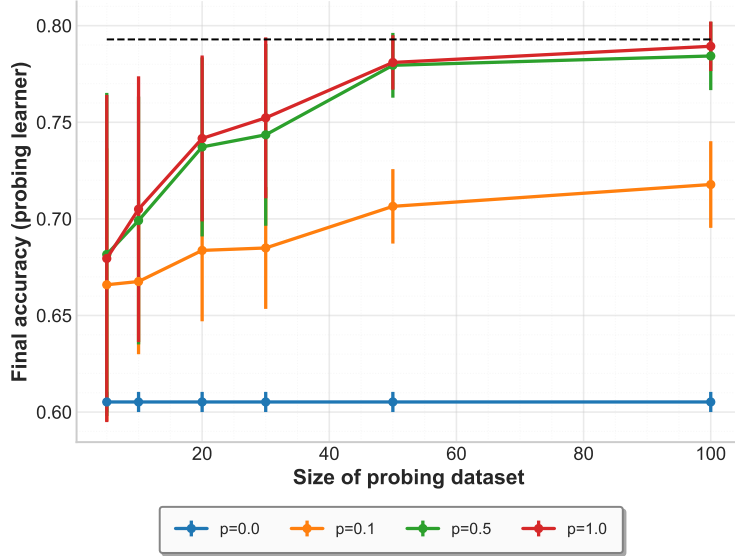


Figure 5: **Performance of probing learner on Census as a function of n .** Error bars show one standard deviation over 10 random seeds.

and variability across runs decreases, consistent with the finite-sample term in Theorem 4, which scales as $1/\sqrt{n}$.

Expt 5: Two-layer neural learners. Figure 6 repeats the Census preference-aware experiments with two-layer ReLU MLP learners. The same qualitative pattern persists: standard MSGD overspecializes, while probing improves the underperforming learner toward the pooled MLP baseline.

7 Discussion

We studied competitive dynamics in machine learning markets where users choose platforms based on inherent preferences and predictive quality. We showed that standard learning dynamics in competitive ML markets converge to overspecialized equilibria, and that peer model probing provably mitigates this failure under identifiable informational conditions. The key insight is that optimizing for observed users creates information barriers that standard dynamics cannot overcome; probing breaks these barriers, with guarantees that degrade gracefully with pseudo-label quality. Our experiments confirm that even small probing datasets suffice to close most of the gap.

User Preference Modeling. Our experiments simulate user preferences via K-means clustering, which provides a controlled environment but may not capture the full complexity of real-world platform choice. Settings with explicit preference data or richer choice models (e.g., multinomial logit with heterogeneous coefficients) merit exploration.

Online Probing. Our analysis considers offline probing, where pseudo-labels are collected once at initialization; this mirrors fixed-teacher distillation workflows and avoids querying peers at different stages of adaptation. Extending to *online probing*, where learners continuously query adapting peers throughout training, introduces co-adaptation dynamics that may lead to instabilities reminiscent of

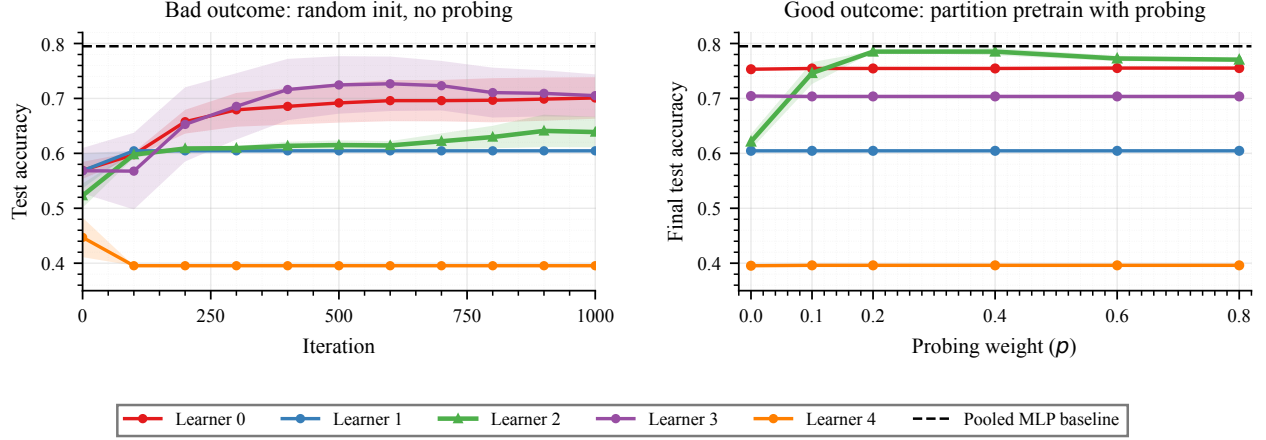


Figure 6: **Census experiments with two-layer neural learners.** We replace each linear logistic learner with a two-layer ReLU MLP ($d \rightarrow 64 \rightarrow 1$) trained by SGD with binary cross-entropy, while keeping the Census train/test split, K-means preference clustering, induced learner rankings, and offline pseudo-label construction fixed. Left: random initialization with standard MSGD and no probing ($\tau = 0.3, p = 0$). Right: partition-pretrained initialization with preference-aware probing by Learner 2 ($\tau = 0.7, \kappa = 0$). Shaded bands show standard error over three seeds.

model collapse [46]. Characterizing when online probing converges—and whether it outperforms offline probing—is an interesting direction for future work.

Convex Losses and Linear Models. Our theoretical analysis is restricted to convex losses (squared and cross-entropy) with linear predictors. The Census MLP experiment in Figure 6 suggests that the qualitative overspecialization and probing-recovery phenomena persist for simple non-convex learners, but extending the convergence and performance guarantees to deep networks remains an important open direction.

References

- [1] Nor Aniza Abdullah, Rasheed Abubakar Rasheed, Mohd Hairul Nizam Nasir, and Md Mujibur Rahman. Eliciting auxiliary information for cold start user recommendation: A survey. *Applied Sciences*, 11(20):9608, 2021.
- [2] Rohan Anil, Gabriel Pereyra, Alexandre Passos, Róbert Ormandi, George E Dahl, and Geoffrey E Hinton. Large scale distributed neural network training through online distillation. In *International Conference on Learning Representations (ICLR)*, 2018.
- [3] Anonymous. Minillm: Knowledge distillation of large language models. *arXiv preprint arXiv:2306.08543*, 2023. URL <https://arxiv.org/abs/2306.08543>.
- [4] Jye E. Beardow. Scroll, click, like, share, repeat: The algorithmic polarisation phenomenon. *ANU Journal of Law & Technology*, 2(1):153–164, 2021. Autumn 2021 issue.
- [5] Omer Ben-Porat and Moshe Tennenholtz. Best response regression. *Advances in Neural Information Processing Systems*, 30, 2017.
- [6] Omer Ben-Porat and Moshe Tennenholtz. Regression equilibrium. In *Proceedings of the 2019 ACM Conference on Economics and Computation*, pages 173–191, 2019.
- [7] James Bennett and Stan Lanning. The Netflix prize. In *Proceedings of KDD Cup and Workshop*. ACM, 2007.
- [8] Vivek S Borkar. *Stochastic approximation: a dynamical systems viewpoint*, volume 9. Springer, 2008.
- [9] Avinandan Bose, Mihaela Curmei, Daniel L Jiang, Jamie Morgenstern, Sarah Dean, Lillian J Ratliff, and Maryam Fazel. Initializing services in interactive ml systems for diverse users. *arXiv preprint arXiv:2312.11846*, 2023.
- [10] Yeshwanth Cherapanamjeri, Constantinos Daskalakis, Andrew Ilyas, and Manolis Zampetakis. What makes a good fisherman? linear regression under self-selection bias. In *Proceedings of the 55th Annual ACM Symposium on Theory of Computing, STOC 2023*. Association for Computing Machinery, 2023. doi: 10.1145/3564246.3585177. URL <https://doi.org/10.1145/3564246.3585177>.
- [11] Wei-Lin Chiang, Zhuohan Li, Zi Lin, Ying Sheng, Zhanghao Wu, Hao Zhang, Lianmin Zheng, Siyuan Zhuang, Yonghao Zhuang, Joseph E. Gonzalez, Ion Stoica, and Eric P. Xing. Vicuna: An open-source chatbot impressing GPT-4 with 90%* ChatGPT quality. <https://lmsys.org/blog/2023-03-30-vicuna/>, March 2023. Accessed: 2025.
- [12] Wei-Lin Chiang, Lianmin Zheng, Ying Sheng, Anastasios Nikolas Angelopoulos, Tianle Li, Dacheng Li, Hao Zhang, Banghua Zhu, Michael Jordan, Joseph E Gonzalez, and Ion Stoica. Chatbot arena: An open platform for evaluating LLMs by human preference. In *Proceedings of the 41st International Conference on Machine Learning*, 2024. URL <https://arxiv.org/abs/2403.04132>.
- [13] Federico Cinus, Marco Minici, Corrado Monti, and Francesco Bonchi. The effect of people recommenders on echo chambers and polarization. In *Proceedings of the Sixteenth International AAAI Conference on Web and Social Media (ICWSM '22)*, pages 90–101. Association for the Advancement of Artificial Intelligence (AAAI), 2022. ICWSM 2022.

- [14] Sarah Dean, Mihaela Curmei, Lillian Ratliff, Jamie Morgenstern, and Maryam Fazel. Emergent specialization from participation dynamics and multi-learner retraining. In *International Conference on Artificial Intelligence and Statistics*, pages 343–351. PMLR, 2024.
- [15] Frances Ding, Moritz Hardt, John Miller, and Ludwig Schmidt. Retiring adult: New datasets for fair machine learning. *Advances in neural information processing systems*, 34:6478–6490, 2021.
- [16] Tony Ginart, Eva Zhang, Yongchan Kwon, and James Zou. Competing ai: How does competition feedback affect machine learning? In *International Conference on Artificial Intelligence and Statistics*, pages 1693–1701. PMLR, 2021.
- [17] António Góis, Mehrnaz Mofakhami, Fernando P. Santos, Gauthier Gidel, and Simon Lacoste-Julien. Performative prediction on games and mechanism design. In *Proceedings of The 28th International Conference on Artificial Intelligence and Statistics*, volume 258 of *Proceedings of Machine Learning Research*, pages 1855–1863. PMLR, 2025.
- [18] Peter M Guadagni and John DC Little. A logit model of brand choice calibrated on scanner data. *Marketing Science*, 2(3):203–238, 1983.
- [19] Arnav Gudibande, Eric Wallace, Charlie Snell, Xinyang Geng, Hao Liu, Pieter Abbeel, Sergey Levine, and Dawn Song. The false promise of imitating proprietary LLMs. In *International Conference on Learning Representations (ICLR)*, 2024.
- [20] Moritz Hardt, Nimrod Megiddo, Christos Papadimitriou, and Mary Wootters. Strategic classification. In *Proceedings of the 2016 ACM Conference on Innovations in Theoretical Computer Science*, pages 111–122. ACM, 2016.
- [21] F Maxwell Harper and Joseph A Konstan. The movielens datasets: History and context. *Acm transactions on interactive intelligent systems (tiis)*, 5(4):1–19, 2015.
- [22] Keegan Harris, Chara Podimata, and Zhiwei Steven Wu. Strategic apple tasting. *Adv. Neural Inf. Process. Syst.*, abs/2306.06250, June 2023.
- [23] Tatsunori Hashimoto, Megha Srivastava, Hongseok Namkoong, and Percy Liang. Fairness without demographics in repeated loss minimization. In *International Conference on Machine Learning*, pages 1929–1938. PMLR, 2018.
- [24] Geoffrey Hinton, Oriol Vinyals, and Jeff Dean. Distilling the knowledge in a neural network. *arXiv preprint arXiv:1503.02531*, 2015. URL <https://arxiv.org/abs/1503.02531>.
- [25] Hong Huang, Bo Zhao, Hao Zhao, Zhou Zhuang, Zhenxuan Wang, Xiaoming Yao, Xinggang Wang, Hai Jin, and Xiaoming Fu. A cross-platform consumer behavior analysis of large-scale mobile shopping data. In *Proceedings of the 2018 World Wide Web Conference, WWW ’18*, pages 1785–1794, Lyon, France, April 2018. International World Wide Web Conferences Steering Committee. doi: 10.1145/3178876.3186169. URL <https://doi.org/10.1145/3178876.3186169>.
- [26] Ruben Interian, Rusl  n G. Marzo, Isela Mendoza, and Celso C. Ribeiro. Network polarization, filter bubbles, and echo chambers: An annotated review of measures and reduction methods. *arXiv preprint*, 2022. arXiv:2207.13799.

- [27] Roozbeh Irani-Kermani, Edward C. Jaenicke, and Ardalan Mirshani. Accommodating heterogeneity in brand loyalty estimation: Application to the U.S. beer retail market. *Journal of Marketing Analytics*, 11(4):820–835, 2023. doi: 10.1057/s41270-022-00187-2. URL <https://doi.org/10.1057/s41270-022-00187-2>.
- [28] Meena Jagadeesan, Michael I Jordan, and Nika Haghtalab. Competition, alignment, and equilibria in digital marketplaces. *Proceedings of the AAAI Conference on Artificial Intelligence*, 37(5):5689–5696, 2023. doi: 10.1609/aaai.v37i5.25706.
- [29] Meena Jagadeesan, Michael I Jordan, Jacob Steinhardt, and Nika Haghtalab. Improved bayes risk can yield reduced social welfare under competition. *arXiv preprint arXiv:2306.14670*, 2023.
- [30] Hailey James, Chirag Nagpal, Katherine A Heller, and Berk Ustun. Participatory personalization in classification. In *Thirty-seventh Conference on Neural Information Processing Systems*, 2023.
- [31] Julie Jiang, Xiang Ren, and Emilio Ferrara. Social media polarization and echo chambers in the context of covid-19: Case study. *JMIRx med*, 2(3):e29570, 2021. doi: 10.2196/29570.
- [32] Yongchan Kwon, Antonio Ginart, and James Zou. Competition over data: how does data purchase affect users? *arXiv preprint arXiv:2201.10774*, 2022.
- [33] Qiang Li, Chung-Yiu Yau, and Hoi-To Wai. Multi-agent performative prediction with greedy deployment and consensus seeking agents. In *Advances in Neural Information Processing Systems*, volume 35, pages 38449–38460, 2022.
- [34] Yixing Li, Yuxian Gu, Li Dong, Dequan Wang, Yu Cheng, and Furu Wei. Direct preference knowledge distillation for large language models. *arXiv preprint arXiv:2406.19774*, 2024. URL <https://arxiv.org/abs/2406.19774>.
- [35] Percy Liang, Rishi Bommasani, Tony Lee, Dimitris Tsipras, Dilara Soylu, Michihiro Yasunaga, Yian Zhang, Deepak Narayanan, Yuhuai Wu, Ananya Kumar, Benjamin Newman, Binhang Yuan, Bobby Yan, Ce Zhang, Christian Cosgrove, Christopher D Manning, Christopher Ré, Diana Acosta-Navas, Drew A Hudson, Eric Zelikman, Esin Durmus, Faisal Ladhak, Frieda Rong, Hongyu Ren, Huaxiu Yao, Jue Wang, Keshav Santhanam, Laurel Orr, Lucia Zheng, Mert Yükeşgönül, Mirac Suzgun, Nathan Kim, Neel Guha, Niladri Chatterji, Omar Khessin, Peter Henderson, Qian Huang, Ryan Chi, Sang Michael Xie, Shibani Santurkar, Surya Ganguli, Tatsunori Hashimoto, Thomas Icard, Tianyi Zhang, Vishrav Chandu, William Wang, Xuechen Xie, Xuechen Zhang, Yushi Wang, Yuntao Zhou, and Yuta Koreeda. Holistic evaluation of language models. *Transactions on Machine Learning Research*, 2023. URL <https://arxiv.org/abs/2211.09110>.
- [36] Eric Mazumdar, Lillian J. Ratliff, and S. Shankar Sastry. On gradient-based learning in continuous games. *SIAM Journal on Mathematics of Data Science*, 2(1):103–131, 2020. doi: 10.1137/18M1231298. URL <https://doi.org/10.1137/18M1231298>.
- [37] John Miller, Juan C Perdomo, and Tijana Zrnic. Outside the echo chamber: Optimizing the performative risk. In *Proceedings of the 38th International Conference on Machine Learning*, volume 139 of *Proceedings of Machine Learning Research*, pages 7710–7720. PMLR, 2021.
- [38] Adhyayan Narang, Evan Faulkner, Dmitriy Drusvyatskiy, Maryam Fazel, and Lillian J Ratliff. Multiplayer performative prediction: Learning in decision-dependent games. *Journal of Machine Learning Research*, 24(202):1–56, 2023.

- [39] Randal S Olson, William La Cava, Patryk Orzechowski, Ryan J Urbanowicz, and Jason H Moore. PMLB: A large benchmark suite for machine learning evaluation and comparison. *BioData Mining*, 10(1):36, 2017.
- [40] Juan Perdomo, Tijana Zrnic, Celestine Mendler-Dünner, and Moritz Hardt. Performative prediction. In *International Conference on Machine Learning*, pages 7599–7609. PMLR, 2020.
- [41] Georgios Piliouras and Fang-Yi Yu. Multi-agent performative prediction: From global stability and optimality to chaos. In *Proceedings of the 24th ACM Conference on Economics and Computation*, pages 1047–1048. ACM, 2023.
- [42] Reilly Raab, Ross Boczar, Maryam Fazel, and Yang Liu. Fair participation via sequential policies. In *AAAI Conference on Artificial Intelligence*, 2024.
- [43] Olga Russakovsky, Jia Deng, Hao Su, Jonathan Krause, Sanjeev Satheesh, Sean Ma, Zhiheng Huang, Andrej Karpathy, Aditya Khosla, Michael Bernstein, Alexander C Berg, and Li Fei-Fei. ImageNet large scale visual recognition challenge. *International Journal of Computer Vision*, 115(3):211–252, 2015.
- [44] H.J. Scudder. Probability of error of some adaptive pattern-recognition machines. *IEEE Transactions on Information Theory*, 11(3):363–371, July 1965.
- [45] Eliot Shekhtman and Sarah Dean. Strategic usage in a multi-learner setting. In *International Conference on Artificial Intelligence and Statistics*. PMLR, 2024.
- [46] Ilia Shumailov, Zakhar Shumaylov, Yiren Zhao, Nicolas Papernot, Ross Anderson, and Yarin Gal. Ai models collapse when trained on recursively generated data. *Nature*, 631(8022):755–759, 2024.
- [47] Jinyan Su and Sarah Dean. Learning from streaming data when users choose. *arXiv [cs.LG]*, June 2024.
- [48] Shicheng Tan, Weng Lam Tam, Yuanchun Wang, Wenwen Gong, Yang Yang, Hongyin Tang, Keqing He, Jiahao Liu, Jingang Wang, Shu Zhao, Peng Zhang, and Jie Tang. Gkd: A general knowledge distillation framework for large-scale pre-trained language model. *arXiv preprint arXiv:2306.06629*, 2023. URL <https://arxiv.org/abs/2306.06629>.
- [49] Rohan Taori, Ishaan Gulrajani, Tianyi Zhang, Yann Dubois, Xuechen Li, Carlos Guestrin, Percy Liang, and Tatsunori B. Hashimoto. Stanford Alpaca: An instruction-following LLaMA model. https://github.com/tatsu-lab/stanford_alpaca, 2023. Accessed: 2025.
- [50] Guanghui Wang, Ioannis Panageas, Georgios Piliouras, and Fang-Yi Yu. Last-iterate convergence for symmetric, general-sum, 2×2 games under the exponential weights dynamic. *arXiv preprint arXiv:2502.08063*, 2025.
- [51] Xiaolu Wang, Chung-Yiu Yau, and Hoi To Wai. Network effects in performative prediction games. In *Proceedings of the 40th International Conference on Machine Learning*, volume 202 of *Proceedings of Machine Learning Research*, pages 36514–36540. PMLR, 2023.
- [52] John Werner. Did deepseek copy off of openai? and what is distillation? *Forbes*, 2025. URL <https://www.forbes.com/sites/johnwerner/2025/01/30/did-deepseek-copy-off-of-openai-and-what-is-distillation/>.

- [53] Xiaohan Xu, Ming Li, Chongyang Tao, Tao Shen, Reynold Cheng, Jinyang Li, Can Xu, Dacheng Tao, and Tianyi Zhou. A survey on knowledge distillation of large language models. *arXiv preprint arXiv:2402.13116*, 2024. URL <https://arxiv.org/abs/2402.13116>.
- [54] Chuanpeng Yang, Wang Lu, Yao Zhu, Yidong Wang, Qian Chen, Chenlong Gao, Bingjie Yan, and Yiqiang Chen. Survey on knowledge distillation for large language models: Methods, evaluation, and application. *arXiv preprint arXiv:2407.01885*, 2024. URL <https://arxiv.org/abs/2407.01885>.
- [55] David Yarowsky. Unsupervised word sense disambiguation rivaling supervised methods. In *33rd Annual Meeting of the Association for Computational Linguistics*, pages 189–196, Cambridge, Massachusetts, USA, June 1995. Association for Computational Linguistics. doi: 10.3115/981658.981684.
- [56] Xueru Zhang, Mohammadmahdi Khaliligarekani, Cem Tekin, et al. Group retention when using machine learning in sequential decision making: the interplay between user dynamics and fairness. *Advances in Neural Information Processing Systems*, 32, 2019.
- [57] Ying Zhang, Tao Xiang, Timothy M Hospedales, and Huchuan Lu. Deep mutual learning. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 4320–4328, 2018.
- [58] Zaiwei Zhu, Runyu Wan, Yingbin Cho, Haoming Luo, Zhuoran Yang, and Zhaoran Wang. Online performative gradient descent for learning nash equilibria in decision-dependent games. In *Advances in Neural Information Processing Systems*, volume 36, 2023.

A Extended Related Work

This section provides detailed comparisons with related work summarized in Section 1.1 of the main paper.

Single-learner user agency. A body of work studies learning dynamics when users are treated as independent agents rather than passive data points. Hashimoto et al. [23] show that empirical risk minimization can cause minority groups to opt out, creating a feedback loop that further degrades minority performance. Zhang et al. [56] study a related dynamic where underrepresented groups receive worse predictions, leading to further data scarcity. James et al. [30] analyze participatory data collection where users choose whether to contribute data. Ben-Porat and Tennenholtz [5] study best-response dynamics in strategic classification. Cherapanamjeri et al. [10] and Harris et al. [22] study settings where data quality or availability depends on the learner’s past performance. We build on this perspective in the multi-learner setting, where inter-learner interactions create additional feedback dynamics.

Multi-learner competition and user choice. Several works study multi-learner user-choice settings with explicitly strategic users who optimize their own utility functions [5, 6, 28, 29, 45]. Ginart et al. [16] provide both empirical and theoretical analysis of how competition drives specialization: their Theorems 4.1–4.3 establish risk ratio bounds showing that competing predictors perform worse on the general population than a single predictor would. Their analysis considers batch retraining dynamics and characterizes the *existence* of performance gaps due to competition. We complement this by analyzing streaming gradient-based dynamics, proving that such dynamics *converge* to overspecialized equilibria (Theorem 2), and proposing probing as a mitigation. Kwon et al. [32] study a related setting where learners may purchase user data, providing primarily empirical analysis of the resulting market dynamics.

Gradient-based dynamics in choice-driven settings. Most closely related to our work, Dean et al. [14] and Su and Dean [47] analyze gradient-based dynamics in choice-driven settings. Su and Dean [47] introduce the MSGD algorithm and prove convergence to stationary points of an aggregate loss across learners. We build directly on their framework: our Algorithm 1 adapts MSGD to our user selection rule (Definition 1), and our potential function (Equation 1) extends theirs. Our Theorem 1 provides an alternate proof of their convergence result using stochastic approximation techniques, which we believe better illuminates the underlying dynamics. Beyond convergence, we analyze generalization to users outside each learner’s observed population—a consideration absent from prior work—and introduce peer probing as a mechanism to restore global competence. Bose et al. [9] propose intelligent initialization schemes to improve outcomes in similar settings, but also focus on the sum of all local losses on the observed distribution.

B Terminology and Preliminary Results

B.1 Dynamical systems terminology

This subsection records standard definitions underlying terms such as “invariant set” and “internally chain transitive” in Lemma 11. Let $\dot{x}(t) = h(x(t))$ be an ODE with locally Lipschitz h , so that solutions are unique. Let $\phi_t(x)$ denote the state at time $t \geq 0$ of the solution initialized at x at time 0.

Definition 4 (Invariant set). A set $A \subseteq \mathbb{R}^d$ is (forward) invariant for the ODE if $\phi_t(x) \in A$ for all $x \in A$ and all $t \geq 0$.

Definition 5 ((ε, T) -chain). Fix $\varepsilon > 0$ and $T > 0$. An (ε, T) -chain in A from x to y is a finite sequence of points $x = x_0, x_1, \dots, x_k = y$ in A and times $t_1, \dots, t_k \geq T$ such that

$$\|\phi_{t_i}(x_{i-1}) - x_i\| < \varepsilon \quad \text{for all } i \in \{1, \dots, k\}.$$

Definition 6 (Internally chain transitive). A compact invariant set A is internally chain transitive if for every $x, y \in A$ and every $\varepsilon > 0$, $T > 0$, there exists an (ε, T) -chain in A from x to y .

B.2 Regularity of Squared and Cross Entropy Losses

Lemma 4 (Regularity of the squared-error loss). Let $\mathcal{X} \subset \mathbb{R}^d$, $\mathcal{Y} \subset \mathbb{R}$, $W \subset \mathbb{R}^d$ be compact. Write

$$\ell(x, y; \theta) = (y - x^\top \theta)^2, \quad B = \max_{\theta \in W} \|\theta\|.$$

Then for all $(x, y) \in \mathcal{X} \times \mathcal{Y}$ and all $\theta_1, \theta_2 \in W$:

(i) ℓ is nonnegative, C^∞ in θ , convex, and β -smooth with $\beta = 2R^2$.

(ii) ℓ is locally Lipschitz on W :

$$|\ell(x, y; \theta_1) - \ell(x, y; \theta_2)| \leq L_W \|\theta_1 - \theta_2\|, \quad L_W = R(2|y| + 2RB).$$

(iii) The Hessian is constant in θ , so $\|\nabla^2 \ell(\theta_1) - \nabla^2 \ell(\theta_2)\| = 0$ (i.e. $\gamma_\ell = 0$).

Proof. (i) One computes $\nabla_\theta \ell = -2x(y - x^\top \theta)$ and $\nabla_\theta^2 \ell = 2xx^\top \succeq 0$. Hence $\ell \geq 0$, is C^∞ , convex, and $\|\nabla \ell(\theta_1) - \nabla \ell(\theta_2)\| \leq 2R^2 \|\theta_1 - \theta_2\|$.

(ii) Observe

$$\ell(\theta_1) - \ell(\theta_2) = (y - x^\top \theta_1)^2 - (y - x^\top \theta_2)^2 = x^\top (\theta_2 - \theta_1) (2y - x^\top (\theta_1 + \theta_2)).$$

Since $\|\theta_i\| \leq B$, $|x^\top (\theta_2 - \theta_1)| \leq R \|\theta_1 - \theta_2\|$ and $|2y - x^\top (\theta_1 + \theta_2)| \leq 2|y| + 2RB$, giving the stated bound.

(iii) Because $\nabla_\theta^2 \ell \equiv 2xx^\top$ is independent of θ , its difference vanishes. □

Lemma 5. Under Assumption 1, for all $z \in \mathcal{Z}$ the loss $\ell(z, \cdot)$ is non-negative, convex, differentiable, locally Lipschitz and β_ℓ -smooth with

$$\beta_\ell = 2R^2 \quad \text{for both the squared loss and cross-entropy loss.}$$

Moreover, let

$$\mathcal{R}(\theta) = \mathbb{E}_{z \sim \mathcal{P}} [\ell(z, \theta)], \quad \theta^* = \arg \min_{\theta} \mathcal{R}(\theta), \quad \epsilon = \mathcal{R}(\theta^*).$$

Then for any θ with $\|\theta - \theta^*\| \leq \gamma$, the following quadratic upper bound holds:

$$\mathcal{R}(\theta) \leq \epsilon + \frac{\beta_\ell}{2} \|\theta - \theta^*\|^2 \leq \epsilon + R^2 \gamma^2.$$

Proof. The first part (non-negativity, convexity, differentiability, local Lipschitzness, and smoothness) follows directly by applying Lemmas 4 and 6 under Assumption 1, which together show each $\ell(z, \theta)$ is β_ℓ -smooth with $\beta_\ell = 2R^2$.

Since $f(\theta) = \mathbb{E}_z[\ell(z, \theta)]$, differentiability and smoothness carry over to f , and we have

$$\|\nabla f(\theta) - \nabla f(\theta^*)\| \leq \beta_\ell \|\theta - \theta^*\|.$$

At the minimizer θ^* , $\nabla f(\theta^*) = 0$. Hence for any θ with $\|\theta - \theta^*\| \leq \gamma$, the gradient satisfies

$$\|\nabla f(\theta)\| \leq \beta_\ell \|\theta - \theta^*\| \leq \beta_\ell \gamma.$$

Applying the standard smoothness inequality (second-order Taylor upper bound) around θ^* ,

$$f(\theta) \leq f(\theta^*) + \langle \nabla f(\theta^*), \theta - \theta^* \rangle + \frac{\beta_\ell}{2} \|\theta - \theta^*\|^2 = \epsilon + \frac{\beta_\ell}{2} \|\theta - \theta^*\|^2,$$

where we used $\nabla f(\theta^*) = 0$ and $f(\theta^*) = \epsilon$. Finally, substituting $\beta_\ell = 2R^2$ gives

$$f(\theta) \leq \epsilon + R^2 \|\theta - \theta^*\|^2 \leq \epsilon + R^2 \gamma^2,$$

completing the proof. $\square$

Lemma 6 (Regularity of the multiclass cross-entropy loss). *Let $\mathcal{X} \subset \mathbb{R}^d$, $\mathcal{Y} = \{1, \dots, K\}$, $W \subset \mathbb{R}^{K \times d}$ compact, and let $\|x\| \leq R$ for all $x \in \mathcal{X}$. Define for $(x, y) \in \mathcal{X} \times \mathcal{Y}$ and $\Theta = (\theta_1, \dots, \theta_K) \in W$*

$$\ell(x, y; \Theta) = -\log \frac{\exp(\theta_y^\top x)}{\sum_{k=1}^K \exp(\theta_k^\top x)}.$$

Then for all x, y, Θ_1, Θ_2 as above:

(i) ℓ is nonnegative, C^∞ in Θ , convex, and β -smooth with $\beta = R^2$.

(ii) ℓ is (locally) Lipschitz on W :

$$|\ell(x, y; \Theta_1) - \ell(x, y; \Theta_2)| \leq L \|\Theta_1 - \Theta_2\|, \quad L = \sqrt{2} R.$$

Proof. The proofs of each subpart are as follows:

(i) Let $p = \text{softmax}(\Theta x) \in \Delta^{K-1}$. One checks

$$\nabla_\Theta \ell = (p - e_y) x^\top, \quad \nabla_\Theta^2 \ell = [\text{diag}(p) - p p^\top] \otimes (x x^\top).$$

Since $\text{diag}(p) - p p^\top \succeq 0$, $\ell \geq 0$, is C^∞ and convex, and

$$\|\nabla_\Theta^2 \ell\| \leq \|x\|^2 \lambda_{\max}(\text{diag}(p) - p p^\top) \leq R^2,$$

it follows that ℓ is β -smooth with $\beta = R^2$.

(ii) By the mean-value theorem,

$$|\ell(\Theta_1) - \ell(\Theta_2)| \leq \sup_{\Theta \in W} \|\nabla_\Theta \ell\| \|\Theta_1 - \Theta_2\|.$$

But

$$\|\nabla_\Theta \ell\| = \|p - e_y\| \|x\| \leq \sqrt{2} R,$$

giving the claimed Lipschitz constant $L = \sqrt{2} R$.

□

Lemma 7. Let $\{h_\theta\}_{\theta \in \Theta}$ be the family of softmax classifiers mapping $x \in \mathbb{R}^d$ to a distribution over C classes,

$$h_\theta(x)_c = \frac{e^{x^\top \theta_c}}{\sum_{j=1}^C e^{x^\top \theta_j}}.$$

Assume:

- (i) $\|x\|_2 \leq R$ for all x in the support of $\mathcal{P}$.
- (ii) There exists $\alpha > 0$ such that for all x, θ, c , we have $h_\theta(x)_c \geq \alpha$.
- (iii) θ, θ^* lie in a convex set $\Theta \subset \mathbb{R}^{C \times d}$.

Then the mapping $\theta \mapsto h_\theta(x)$ is Lipschitz in the ℓ_1 -norm: for all $\theta, \theta^* \in \Theta$,

$$\|h_\theta(x) - h_{\theta^*}(x)\|_1 \leq L_h \|\theta - \theta^*\|_2, \quad L_h = R\sqrt{2(1-\alpha)}.$$

Proof. By the multivariate mean-value theorem, there exists $\bar{\theta}$ on the line segment between θ^* and θ such that

$$h_\theta(x) - h_{\theta^*}(x) = D_\theta h(\bar{\theta}; x) (\theta - \theta^*),$$

where $D_\theta h(\bar{\theta}; x) \in \mathbb{R}^{C \times (Cd)}$ is the Jacobian w.r.t. all parameters. Write its c th row as the gradient vector $\nabla_\theta h_c(\bar{\theta}; x)$. A direct calculation shows for each class c and each block k :

$$\nabla_{\theta_k} h_c(\bar{\theta}; x) = h_c(\mathbf{1}\{c = k\} - h_k) x \implies \|\nabla_\theta h_c(\bar{\theta}; x)\|_2 \leq \|x\|_2 h_c \sqrt{(1 - h_c)^2 + \sum_{k \neq c} h_k^2}.$$

Using $\sum_k h_k = 1$, $h_k \leq 1$ for each k , and the lower-bound $h_c \geq \alpha$, one checks

$$(1 - h_c)^2 + \sum_{k \neq c} h_k^2 \leq (1 - h_c)^2 + (1 - h_c) = (1 - h_c)(2 - h_c) \leq 2(1 - h_c) \leq 2(1 - \alpha),$$

where the first inequality uses $\sum_{k \neq c} h_k^2 \leq \sum_{k \neq c} h_k = 1 - h_c$. Hence

$$\|\nabla_\theta h_c(\bar{\theta}; x)\|_2 \leq R h_c \sqrt{2(1 - \alpha)}.$$

Therefore the operator norm from ℓ_2 to ℓ_1 satisfies

$$\|D_\theta h(\bar{\theta}; x)\|_{2 \rightarrow 1} = \sum_{c=1}^C \|\nabla_\theta h_c(\bar{\theta}; x)\|_2 \leq R\sqrt{2(1 - \alpha)} \sum_{c=1}^C h_c = R\sqrt{2(1 - \alpha)}.$$

Putting these together gives

$$\|h_\theta(x) - h_{\theta^*}(x)\|_1 \leq \|D_\theta h(\bar{\theta}; x)\|_{2 \rightarrow 1} \|\theta - \theta^*\|_2 \leq L_h \|\theta - \theta^*\|_2,$$

as claimed. □

Lemma 8. Under Assumption 1, for all $z \in \mathcal{Z}$, the loss $\ell(z, \cdot)$ is non-negative, convex, differentiable, locally Lipschitz, and β_ℓ -smooth.

Proof. The result follows by applying Lemmas 4 and 6 with $|x| \leq R$ and choosing $\beta_\ell = \max 2R^2, R^2 = 2R^2$. Each loss is nonnegative, convex, differentiable, locally Lipschitz, and smooth under these bounds. □

C Bad Outcome

Notation Recall the potential function from the main text:

$$f(\Theta) = \sum_{i=1}^m \left(\tau \alpha_i \mathbb{E}_{z \sim \mathcal{P}_i} [\ell(z, \theta_i)] + (1 - \tau) a_i(\Theta) \mathbb{E}_{z \sim \mathcal{D}_i(\Theta)} [\ell(z, \theta_i)] \right) \quad (4)$$

From Lemma 1, we have that the MSGD iterates converge almost surely to the set $\{\nabla f(\Theta) = 0\}$. Additionally, for learner $i \in [m]$, define $\bar{\theta}_i$ as the solution that each learner would choose if they trained independently on the subpopulation of users who rank them highest:

$$\bar{\theta}_i = \arg \min_{\theta_i} \mathbb{E}_{x \sim \mathcal{P}_i} \ell(x, \theta_i) \quad (5)$$

Example 1. Specify the family of instances $\mathcal{G}_{bad}(\tau, C)$ as follows.

1. *Distribution $\mathcal{P}$:* is defined as a mixture of subpopulations. $\mathcal{P} = \alpha \mathcal{P}_1 + (1 - \alpha) \mathcal{P}_2$. Here, $\{\mathcal{P}_1, \mathcal{P}_2\}$ are 2 subpopulations. For either subpopulation, the covariates are generated from the zero mean and unit-variance uniform distribution:

$$x \sim \text{Unif}[-\sqrt{3}, \sqrt{3}] \quad \text{for } (x, y) \sim \mathcal{P}_i, \quad (6)$$

For each subpopulation, the response variable is generated as:

$$y = Cx \quad \text{for } (x, y) \sim \mathcal{P}_1 \quad (7)$$

$$y = -x \quad \text{for } (x, y) \sim \mathcal{P}_2 \quad (8)$$

2. *Loss function:* $\ell(x, y, \theta) = (y - \theta^T x)^2$ is the squared loss
3. *Ranking* $\pi(z) = i$ for $(x, y) \sim \mathcal{P}_i$.

Lemma 9. Consider the 1-D bad-outcome family in Example 1 with mixture weight $\alpha \in (0, 1)$ and slope parameter $C > 1$, and let $\mathcal{R}(\theta) = \mathbb{E}_{(x, y) \sim \mathcal{P}} [(y - \theta x)^2]$.

- (i) The least-squares predictor on the full mixture, $\theta^* = \alpha C - (1 - \alpha)$, satisfies

$$\mathcal{R}(\theta^*) = \alpha(1 - \alpha)(C + 1)^2.$$

- (ii) The specialist trained on $\mathcal{P}_1$ is $\bar{\theta}_1 = C$, and its mixture risk is

$$\mathcal{R}(\bar{\theta}_1) = (1 - \alpha)(C + 1)^2.$$

- (iii) The specialist trained on $\mathcal{P}_2$ is $\bar{\theta}_2 = -1$, and its mixture risk is

$$\mathcal{R}(\bar{\theta}_2) = \alpha(C + 1)^2.$$

Proof. We have that $\mathbb{E}[x] = 0$ and $\mathbb{E}[x^2] = 1$, so that the squared risk reduces to $(\beta - \theta)^2$ when the true slope is β .

(i) Global compromise. On the mixture, the conditional label is linear: $y = \beta x$ with $\beta = \alpha C + (1 - \alpha)(-1) = \alpha C - (1 - \alpha)$. For squared loss with $\mathbb{E}[x] = 0$ and $\mathbb{E}[x^2] = 1$, the population least-squares minimizer is the regression slope, so $\theta^* = \beta$. The corresponding risk is

$$\mathcal{R}(\theta^*) = \alpha(C - \beta)^2 + (1 - \alpha)(-1 - \beta)^2 = \alpha(1 - \alpha)(C + 1)^2.$$

(ii) **Specialist for $\mathcal{P}_1$.** Training least squares on $\mathcal{P}_1$ alone yields $\bar{\theta}_1 = C$ since $\mathbb{E}[xy] = C$ and $\mathbb{E}[x^2] = 1$. The mixture risk is

$$\mathcal{R}(\bar{\theta}_1) = \alpha(C - C)^2 + (1 - \alpha)(-1 - C)^2 = (1 - \alpha)(C + 1)^2.$$

(iii) **Specialist for $\mathcal{P}_2$.** Training least squares on $\mathcal{P}_2$ gives $\bar{\theta}_2 = -1$ (its true slope). The mixture risk is

$$\mathcal{R}(\bar{\theta}_2) = \alpha(C + 1)^2 + (1 - \alpha)(-1 - (-1))^2 = \alpha(C + 1)^2.$$

□

Lemma 10. Let $(\tilde{\theta}_1, \tilde{\theta}_2) \in \{\nabla f(\Theta) = 0\}$ be a stationary point of $f(\Theta)$ (Equation 1). Then, under the assumption that $\tau \geq \frac{1}{2}$, it must be true that $\mathcal{D}_1(\tilde{\theta}_1, \tilde{\theta}_2) = \mathcal{P}_1$ and $\mathcal{D}_2(\tilde{\theta}_1, \tilde{\theta}_2) = \mathcal{P}_2$.

Proof. The proof proceeds by considering the decisions of the loss-minimizing users, and shows that the stationary point condition implies the conditions on $\mathcal{D}_i(\Theta)$ in the lemma statement. Consider any (x, y) from $\mathcal{P}_1$. For any $\theta \in \mathbb{R}$, we have that

$$\ell(x, y, \theta) = (Cx - \theta x)^2 = (C - \theta)^2 x^2. \quad (9)$$

Hence,

$$\arg \min_{i \in \{1, 2\}} \ell(x, y, \theta_i) = \arg \min_{i \in \{1, 2\}} (C - \theta_i)^2 x^2 = \arg \min_{i \in \{1, 2\}} |C - \theta_i|.$$

Thus every loss-minimizing user in $\mathcal{P}_1$ picks the learner whose parameter is closer to C :

$$M(z; \Theta) = \arg \min_{i \in \{1, 2\}} |C - \theta_i|, \quad z \in \mathcal{P}_1.$$

Likewise, every loss-minimizing user in $\mathcal{P}_2$ picks the learner whose parameter is closer to -1 :

$$M(z; \Theta) = \arg \min_{i \in \{1, 2\}} |-1 - \theta_i|, \quad z \in \mathcal{P}_2.$$

Hence, within each subpopulation $\mathcal{P}_i$, all loss-minimizing users make the same choice. Consequently, for any Θ , the observed-data distributions $\mathcal{D}_i(\Theta)$ can only take one of the following four forms:

1. $\mathcal{D}_1(\Theta) = \mathcal{P}$ and $\mathcal{D}_2(\Theta)$ has zero mass under $\mathcal{P}$
2. $\mathcal{D}_1(\Theta)$ has zero mass under $\mathcal{P}$ and $\mathcal{D}_2(\Theta) = \mathcal{P}$
3. $\mathcal{D}_1(\Theta) = \mathcal{P}_2$ and $\mathcal{D}_2(\Theta) = \mathcal{P}_1$
4. $\mathcal{D}_1(\Theta) = \mathcal{P}_1$ and $\mathcal{D}_2(\Theta) = \mathcal{P}_2$

For any stationary point $(\tilde{\theta}_1, \tilde{\theta}_2)$, we consider each case separately and show that all cases other than the last one lead to a contradiction.

Case 1: In this case, the second learner observes only negative labels, and the first observes a mix of both positive and negative labels. Hence, it should be intuitively true that $\tilde{\theta}_1 > \tilde{\theta}_2 > -1$. This would lead to a contradiction since this would imply that users from $\mathcal{P}_2$ would strictly prefer the second learner over the first one. Now, we can verify this formally. Define the total probability masses seen by each learner

$$M_1 = \tau \alpha + (1 - \tau) \alpha + (1 - \tau)(1 - \alpha) = \alpha + (1 - \tau)(1 - \alpha), \quad (10)$$

$$M_2 = \tau(1 - \alpha). \quad (11)$$

The stationary conditions are given by

$$\tilde{\theta}_1 = \arg \min_{\theta_1} \tau \alpha \mathbb{E}_{P_1} [(y - \theta_1 x)^2] + (1 - \tau) \alpha \mathbb{E}_{P_1} [(y - \theta_1 x)^2] + (1 - \tau)(1 - \alpha) \mathbb{E}_{P_2} [(y - \theta_1 x)^2], \quad (12)$$

$$\tilde{\theta}_2 = \arg \min_{\theta_2} \tau (1 - \alpha) \mathbb{E}_{P_2} [(y - \theta_2 x)^2]. \quad (13)$$

We can solve these optimization problems in closed form:

$$\tilde{\theta}_1 = \frac{\alpha C - (1 - \tau)(1 - \alpha)}{M_1}, \quad (14)$$

$$\tilde{\theta}_2 = -1. \quad (15)$$

Now,

$$\tilde{\theta}_1 - \tilde{\theta}_2 = \tilde{\theta}_1 + 1 = \frac{\alpha(C + 1)}{M_1} > 0,$$

which holds for every $\tau \in [0, 1]$ since $\alpha > 0$ and $C > 0$. Since $\tilde{\theta}_2 = -1$ and $\tilde{\theta}_1 > -1$, $\mathcal{P}_2$ users would strictly prefer learner 2 over learner 1, leading to a contradiction.

Case 2: In this case, the first learner observes only positive labels, and the second observes a mix of both positive and negative labels. Hence, it is easy to verify using a similar procedure as the previous case that $C > \tilde{\theta}_1 > \tilde{\theta}_2$, which would lead to a contradiction since $\mathcal{P}_1$ users would strictly prefer learner 1.

Case 3: We follow a similar argument as in Case 1. Define the total probability masses seen by each learner

$$M_1 = \tau \alpha + (1 - \tau)(1 - \alpha), \quad (16)$$

$$M_2 = \tau(1 - \alpha) + (1 - \tau)\alpha. \quad (17)$$

The stationary points in this case are given by

$$\tilde{\theta}_1 = \frac{\tau \alpha C - (1 - \tau)(1 - \alpha)}{\tau \alpha + (1 - \tau)(1 - \alpha)}, \quad \tilde{\theta}_2 = \frac{(1 - \tau) \alpha C - \tau(1 - \alpha)}{(1 - \tau) \alpha + \tau(1 - \alpha)}.$$

It is easy to verify that both $\tilde{\theta}_1$ and $\tilde{\theta}_2$ are greater than -1 . Hence, the distance from -1 is given by

$$d_1 = |\tilde{\theta}_1 - (-1)| = \frac{\tau \alpha (C + 1)}{\tau \alpha + (1 - \tau)(1 - \alpha)},$$

$$d_2 = |\tilde{\theta}_2 - (-1)| = \frac{(1 - \tau) \alpha (C + 1)}{(1 - \tau) \alpha + \tau(1 - \alpha)}.$$

A straightforward comparison gives

$$d_1 > d_2 \iff \tau M_2 > (1 - \tau) M_1 \iff (1 - \alpha)(2\tau - 1) > 0 \iff \tau > \frac{1}{2}.$$

Therefore, whenever $\tau > \frac{1}{2}$, the free users from $\mathcal{P}_2$ strictly prefer learner 2 over learner 1. This contradicts the Case 3 hypothesis that $\mathcal{D}_1(\Theta) = \mathcal{P}_2$. Hence Case 3 cannot be a stationary point when $\tau > 1/2$. Having ruled out Cases 1–3, only the fourth configuration remains, completing the proof. $\square$

Theorem 2. *Let Assumptions 1–4 hold. For any $\tau \geq \frac{1}{2}$ and any choice of ϵ, Γ with $0 < \epsilon < \Gamma$, there exists an instance $\mathcal{G}$ such that:*

1. *There exists θ^* with $\mathcal{R}(\theta^*) \leq \epsilon$.*
2. *The MSGD iterates converge to a unique stationary point $\bar{\Theta}$ where $\mathcal{R}(\bar{\theta}_i) \geq \Gamma$ for some learner $i \in [m]$.*

Proof. The proof follows by considering Example 1.

Part (i) The proof follows from Lemma 9. Choosing $\alpha = \frac{\epsilon}{(C+1)^2}$ satisfies the condition.

Characterizing the stationary points From Lemma 1, we have that the MSGD iterates converge almost surely to the set $\{\nabla f(\Theta) = 0\}$. By Lemma 10, we have that $(\tilde{\theta}_1, \tilde{\theta}_2) \in \{\nabla f(\Theta) = 0\}$ if and only if $\mathcal{D}_1(\tilde{\theta}_1, \tilde{\theta}_2) = \mathcal{P}_1$ and $\mathcal{D}_2(\tilde{\theta}_1, \tilde{\theta}_2) = \mathcal{P}_2$. Hence, any stationary point must satisfy

$$\tilde{\theta}_1 = \arg \min_{\theta \in \mathbb{R}} \alpha \mathbb{E}_{z \sim \mathcal{P}_1} [\ell(z, \theta)]$$

and

$$\tilde{\theta}_2 = \arg \min_{\theta \in \mathbb{R}} (1 - \alpha) \mathbb{E}_{z \sim \mathcal{P}_2} [\ell(z, \theta)]$$

Clearly, $(\bar{\theta}_1, \bar{\theta}_2)$ is the unique solution to these equations, and MSGD must converge to this point almost surely.

Characterizing the loss From Lemma 9, at the stationary point we have $\bar{\theta}_1 = C$ and $\bar{\theta}_2 = -1$, so the mixture risks are

$$\mathcal{R}(\bar{\theta}_1) = (1 - \alpha)(C + 1)^2, \quad \mathcal{R}(\bar{\theta}_2) = \alpha(C + 1)^2.$$

Moreover, the optimal global (compromise) risk satisfies $\mathcal{R}(\theta^*) \leq \epsilon$ by our choice in Part (i). To ensure learner 1's risk exceeds Γ , choose

$$C = \sqrt{\Gamma + \epsilon} - 1, \quad \alpha = \frac{\epsilon}{(C + 1)^2}.$$

Then $\alpha \in (0, 1)$ and

$$\mathcal{R}(\theta^*) = \alpha(1 - \alpha)(C + 1)^2 \leq \alpha(C + 1)^2 = \epsilon,$$

while

$$\mathcal{R}(\bar{\theta}_1) = (1 - \alpha)(C + 1)^2 = (C + 1)^2 - \epsilon \geq \Gamma.$$

Thus the two claims in the theorem are simultaneously satisfied. $\square$

D Convergence of Algorithm 1 and Algorithm 2

D.1 Convergence of Algorithm 1

The convergence results for MSGD (Algorithm 1) follow as special cases of the MSGD-P analysis when the probing set $U = \emptyset$. In this case, all probing terms vanish and the augmented potential $\tilde{f}$ reduces to the original potential f .

Lemma 2. *Let Assumptions 1–4 hold. Then the iterates $\{\Theta^t\}$ of Algorithm 1 converge to a compact connected internally chain transitive invariant set of the ODE $\dot{\Theta} = -\nabla f(\Theta)$.*

Algorithm 3 Multi-learner Streaming Gradient Descent

Require: loss function $\ell(\cdot, \cdot) \geq 0$; Initial models $\Theta^0 = (\theta_1^0, \dots, \theta_m^0)$; Learning rate $\{\eta^t\}_{t=1}^{T+1}$

- 1: **for** $t = 0, 1, 2, \dots, T$ **do**
- 2: Sample data point $z^t \sim \mathcal{P}$
- 3: User selects model $i = M(z^t; \Theta^t)$
- 4: $\theta_i^{t+1} \leftarrow \theta_i^t - \eta^t \nabla \ell(z^t, \theta_i^t)$
- 5: **end for**
- 6: **return** Θ^T

Proof. This follows from the proof of Theorem 3 with $U = \emptyset$. When $U = \emptyset$:

- The indicator $\mathbf{1}_{i \in U} = 0$ for all $i \in [m]$, so all probing terms vanish.
- The augmented potential reduces to $\tilde{f}(\Theta) = f(\Theta)$.
- The ODE (20) becomes $\dot{\Theta} = -\nabla f(\Theta)$.

The stochastic approximation argument proceeds identically: by Lemma 11, verifying local Lipschitzness of $-\nabla f$ (Lemma 8, Lemma 14), the step-size conditions (Assumption 2), the martingale variance bound (Lemma 12 with $U = \emptyset$, so $K_{\text{probe}} = 0$), and bounded iterates (Assumption 4), the iterates converge almost surely to a compact connected internally chain transitive invariant set of $\dot{\Theta} = -\nabla f(\Theta)$. $\square$

Theorem 1. *Let Assumptions 1–4 hold. Then the iterates $\{\Theta^t\}$ of Algorithm 1 converge to the set of stationary points $\{\Theta : \nabla f(\Theta) = 0\}$ almost surely.*

Proof. Follows directly from Theorem 3 with $U = \emptyset$. $\square$

D.2 Convergence of MSGD-P (Algorithm 2)

Define the empirical probing loss and augmented potential:

$$\hat{L}_i(\theta_i) = \frac{1}{n} \sum_{q=1}^n \ell(\tilde{z}_i^q, \theta_i), \quad i \in U, \quad (18)$$

$$\tilde{f}(\Theta) = f(\Theta) + p \sum_{i \in U} \left(\hat{L}_i(\theta_i) + \frac{\lambda}{2} \|\theta_i\|^2 \right). \quad (19)$$

Define the ordinary differential equation (ODE) $\dot{\Theta} = \tilde{F}(\Theta)$ with:

$$\tilde{F}_i(\Theta) = - \left(\tau \alpha_i \mathbb{E}_{z \sim \mathcal{P}_i} [\nabla \ell(z, \theta_i)] + (1 - \tau) a_i(\Theta) \mathbb{E}_{z \sim \mathcal{D}_i(\Theta)} [\nabla \ell(z, \theta_i)] + p \mathbf{1}_{i \in U} (\nabla \hat{L}_i(\theta_i) + \lambda \theta_i) \right) \quad (20)$$

Theorem 3. *Let Assumptions 1–3, Assumption 4 (as applied to $\tilde{f}$), and Assumption 5 hold. Then, the iterates $\{\Theta^t\}$ of Algorithm 2 converge to the set of stationary points $\{\Theta : \nabla \tilde{f}(\Theta) = 0\}$ almost surely.*

Proof. The proof follows the stochastic approximation template by showing the iterates track the ODE $\dot{\Theta} = \tilde{F}(\Theta)$ and then using a Lyapunov argument for $\tilde{f}$.

Let $z_1 \sim \mathcal{P}$, $z_2 \sim \mathcal{P}_i$, and $z_3 \sim \mathcal{D}_i(\Theta^t)$. For the on-platform update, define for each coordinate i the random direction

$$g_{i,\text{plat}}^t(\Theta^t) = \begin{cases} \nabla \ell(z_2, \theta_i^t), & \text{w.p. } \tau \alpha_i, \\ \nabla \ell(z_3, \theta_i^t), & \text{w.p. } (1 - \tau) a_i(\Theta^t), \\ 0, & \text{otherwise.} \end{cases}$$

For the probing update, for each $j \in U$ sample $\tilde{z}_j^t$ uniformly from the fixed dataset $\mathfrak{D}_j$ and set

$$g_{i,\text{probe}}^t(\Theta^t) = p \mathbf{1}_{i \in U} (\nabla \ell(\tilde{z}_i^t, \theta_i^t) + \lambda \theta_i^t).$$

Let $\tilde{g}_i^t(\Theta^t) = g_{i,\text{plat}}^t(\Theta^t) + g_{i,\text{probe}}^t(\Theta^t)$ and $\tilde{g}^t = (\tilde{g}_1^t, \dots, \tilde{g}_m^t)$. The algorithmic iterate satisfies

$$\Theta^{t+1} = \Theta^t - \eta_t \tilde{g}^t(\Theta^t) = \Theta^t - \eta_t (\tilde{F}(\Theta^t) + v^t),$$

where $v^t = \tilde{g}^t(\Theta^t) - \mathbb{E}[\tilde{g}^t(\Theta^t) \mid \mathcal{F}_t]$.

Drift identity: By construction and by uniform sampling from $\mathfrak{D}_i$,

$$\begin{aligned} \mathbb{E}[\tilde{g}_i^t(\Theta^t) \mid \mathcal{F}_t] &= \tau \alpha_i \mathbb{E}_{z \sim \mathcal{P}_i} [\nabla \ell(z, \theta_i^t)] + (1 - \tau) a_i(\Theta^t) \mathbb{E}_{z \sim \mathcal{D}_i(\Theta^t)} [\nabla \ell(z, \theta_i^t)] \\ &\quad + p \mathbf{1}_{i \in U} (\nabla \hat{L}_i(\theta_i^t) + \lambda \theta_i^t) \\ &= -\tilde{F}_i(\Theta^t). \end{aligned}$$

Below, we verify the assumptions of Lemma 11:

- From Lemma 8 and Lemma 14, the drift $\tilde{F}(\Theta)$ is locally Lipschitz; the probe part $\theta_i \mapsto p(\nabla \hat{L}_i(\theta_i) + \lambda \theta_i)$ is a finite average of locally Lipschitz gradients plus a globally Lipschitz linear term.
- From Assumption 2, the step sizes satisfy $\sum_t \eta_t = \infty$ and $\sum_t \eta_t^2 < \infty$.
- From Lemma 12, augmented to include the probe term, there exists $K > 0$ such that $\mathbb{E}[\|v^t\|^2 \mid \mathcal{F}_t] \leq K(1 + \|\Theta^t\|^2)$.
- From Assumption 4, we have $\sup_t \|\Theta^t\| < \infty$ almost surely.

By Lemma 11, the iterates converge almost surely to a compact, connected, internally chain transitive invariant set of $\dot{\Theta} = \tilde{F}(\Theta)$. Since $\tilde{F}(\Theta) = -\nabla \tilde{f}(\Theta)$, along ODE trajectories

$$\frac{d}{dt} \tilde{f}(\Theta(t)) = \langle \nabla \tilde{f}, \dot{\Theta} \rangle = -\|\nabla \tilde{f}\|^2 \leq 0,$$

with equality iff $\nabla \tilde{f} = 0$. Thus $\tilde{f}$ is a strict Lyapunov function and the only invariant sets are stationary points, yielding the claim. $\square$

Lemma 11 (Borkar [8], Chapter 2, Theorem 2). *Let $\{x_n\}$ be a sequence generated by the stochastic approximation algorithm*

$$x_{n+1} = x_n + a(n)[h(x_n) + M_{n+1}], \quad n \geq 0,$$

where:

1. $h : \mathbb{R}^d \rightarrow \mathbb{R}^d$ is Lipschitz continuous
2. The step sizes $\{a(n)\}$ satisfy $\sum_n a(n) = \infty$ and $\sum_n a(n)^2 < \infty$

3. $\{M_n\}$ is a martingale difference sequence satisfying $E[||M_{n+1}||^2 | \mathcal{F}_n] \leq K(1 + ||x_n||^2)$ for some $K > 0$

4. $\sup_n ||x_n|| < \infty$ almost surely

Then almost surely, the sequence $\{x_n\}$ converges to a (possibly sample path dependent) compact connected internally chain transitive invariant set of the ODE

$$\dot{x}(t) = h(x(t)).$$

Lemma 12 (Martingale variance bound). *Suppose that Assumption 1 holds and that the probing datasets $\{\mathfrak{D}_i\}_{i \in U}$ are fixed finite sets. Let $g^t(\Theta^t) = (g_1^t(\Theta^t), \dots, g_m^t(\Theta^t))$ be the stochastic update vector defined by*

$$g_i^t(\Theta^t) = g_{i,\text{plat}}^t(\Theta^t) + g_{i,\text{probe}}^t(\Theta^t),$$

where

$$g_{i,\text{plat}}^t(\Theta^t) = \begin{cases} \nabla \ell(z^t, \theta_i^t), & \text{if } i = i_t, \\ 0, & \text{if } i \neq i_t, \end{cases} \quad g_{i,\text{probe}}^t(\Theta^t) = p \mathbf{1}_{i \in U} (\nabla \ell(\tilde{z}_i^t, \theta_i^t) + \lambda \theta_i^t),$$

with $\tilde{z}_i^t$ sampled uniformly from $\mathfrak{D}_i$ independently of i_t . Define the filtration $\mathcal{F}_t = \sigma(\Theta^0, z^1, \dots, z^{t-1}, \{\mathfrak{D}_i\}_{i \in U})$, which contains all information available prior to time t . Define the martingale difference sequence:

$$v^t = g^t(\Theta^t) - \mathbb{E}[g^t(\Theta^t) | \mathcal{F}_t].$$

Then there exists a constant $K > 0$ (made explicit below) such that:

$$\mathbb{E}[||v^t||^2 | \mathcal{F}_t] \leq K (1 + ||\Theta^t||^2).$$

Proof. Gradient bounds. From Lemmas 4 and 6, for all (x, y) with $||x|| \leq R$ and all θ we have bounds of the form

$$||\nabla_\theta \ell(x, y; \theta)|| \leq A_0 + A_1 ||\theta||,$$

with

$$(A_0, A_1) = \begin{cases} (2R Y_{\max}, 2R^2), & \text{(squared loss),} \\ (\sqrt{2} R, 0), & \text{(cross-entropy).} \end{cases}$$

For the probe term under squared loss, labels are pseudo-labels $\hat{y} = \text{median}\{\langle x, \theta_0^j \rangle : j \in [m]\}$ so that $|\hat{y}| \leq R \max_{j \in [m]} ||\theta_0^j|| =: Y_{\max, \text{probe}}$, giving $||\nabla_\theta \ell(x, \hat{y}; \theta)|| \leq \tilde{A}_0 + A_1 ||\theta||$ with $\tilde{A}_0 := 2R Y_{\max, \text{probe}}$. Including the L2 regularization term (which appears only in the probe update) and using the triangle inequality,

$$||p(\nabla_\theta \ell(x, \hat{y}; \theta) + \lambda \theta)|| \leq p \tilde{A}_0 + p(A_1 + \lambda) ||\theta||.$$

Define the block constants

$$K_{\text{plat}} := 2 \max(A_0^2, A_1^2), \quad K_{\text{probe}} := 2 \max((p \tilde{A}_0)^2, (p(A_1 + \lambda))^2),$$

where for cross-entropy we take $A_1 = 0$ and use the same $A_0 = \sqrt{2}R$ for both platform and probe.

Variance decomposition. By definition of v^t and $\text{Var}(Y) = \mathbb{E}[||Y||^2] - \|\mathbb{E}[Y]\|^2$,

$$\mathbb{E}[||v^t||^2 | \mathcal{F}_t] \leq \mathbb{E}[||g^t(\Theta^t)||^2 | \mathcal{F}_t].$$

We bound the RHS. Since g^t has at most one nonzero platform block and up to $|U|$ nonzero probe blocks,

$$\|g^t(\Theta^t)\|^2 = \sum_{i=1}^m \|g_{i,\text{plat}}^t + g_{i,\text{probe}}^t\|^2 \leq 2 \sum_{i=1}^m \|g_{i,\text{plat}}^t\|^2 + 2 \sum_{i=1}^m \|g_{i,\text{probe}}^t\|^2.$$

Taking conditional expectations and using the selection probabilities for the platform block as in the MSGD lemma,

$$\begin{aligned} \mathbb{E}[\|g^t(\Theta^t)\|^2 \mid \mathcal{F}_t] &\leq 2 \sum_{i=1}^m \mathbb{P}(i_t = i \mid \mathcal{F}_t) \mathbb{E}[\|\nabla \ell(z^t, \theta_i^t)\|^2 \mid \mathcal{F}_t, i_t = i] \\ &\quad + 2 \sum_{j \in U} \mathbb{E}[\|p \nabla \ell(\tilde{z}_j^t, \theta_j^t)\|^2 \mid \mathcal{F}_t]. \end{aligned}$$

Applying the block bounds gives

$$\begin{aligned} \mathbb{E}[\|g^t(\Theta^t)\|^2 \mid \mathcal{F}_t] &\leq 2K_{\text{plat}} \sum_{i=1}^m \mathbb{P}(i_t = i \mid \mathcal{F}_t) (1 + \|\theta_i^t\|^2) \\ &\quad + 2K_{\text{probe}} \sum_{j \in U} (1 + \|\theta_j^t\|^2) \\ &\leq 2K_{\text{plat}} (1 + \|\Theta^t\|^2) + 2|U|K_{\text{probe}} + 2K_{\text{probe}} \|\Theta^t\|^2 \\ &= 2(K_{\text{plat}} + |U|K_{\text{probe}}) + 2(K_{\text{plat}} + K_{\text{probe}}) \|\Theta^t\|^2. \end{aligned}$$

Therefore, setting

$$K := 2 \max\{K_{\text{plat}} + |U|K_{\text{probe}}, K_{\text{plat}} + K_{\text{probe}}\}$$

yields

$$\mathbb{E}[\|v^t\|^2 \mid \mathcal{F}_t] \leq K (1 + \|\Theta^t\|^2).$$

This completes the proof. $\square$

Lemma 13 (Gradient of the augmented objective). *For every $i \in [m]$ the gradient of $\tilde{f}$ with respect to θ_i is*

$$\begin{aligned} \nabla_{\theta_i} \tilde{f}(\Theta) &= \tau \alpha_i \mathbb{E}_{z \sim \mathcal{P}_i} [\nabla \ell(z, \theta_i)] + (1 - \tau) a_i(\Theta) \mathbb{E}_{z \sim \mathcal{D}_i(\Theta)} [\nabla \ell(z, \theta_i)] \\ &\quad + p \mathbf{1}_{i \in U} (\nabla \hat{L}_i(\theta_i) + \lambda \theta_i). \end{aligned}$$

Proof. We treat a single index i ; all other coordinates of Θ are held fixed. The term $(1 - \tau) a_i(\Theta) \mathbb{E}_{z \sim \mathcal{D}_i(\Theta)} [\ell(z, \theta_i)]$ depends on θ_i both explicitly (inside ℓ) and implicitly through $a_i(\Theta)$ and $\mathcal{D}_i(\Theta)$. Lemma 4.3 of Su and Dean [47] proves that for any differentiable ℓ satisfying the stated regularity,

$$\nabla_{\theta_i} \left(a_i(\Theta) \mathbb{E}_{z \sim \mathcal{D}_i(\Theta)} [\ell(z, \theta_i)] \right) = a_i(\Theta) \mathbb{E}_{z \sim \mathcal{D}_i(\Theta)} [\nabla \ell(z, \theta_i)].$$

Multiplying by $(1 - \tau)$ yields the corresponding contribution, while the $\tau \alpha_i$ -term is immediate. Local L-Lipschitzness of ℓ ensures the required directional limits.

The probing part depends only on θ_i through the finite average $\hat{L}_i(\theta_i) = \frac{1}{n} \sum_{q=1}^n \ell(\tilde{z}_i^q, \theta_i)$ and the regularization term $\frac{\lambda}{2} \|\theta_i\|^2$, so $\nabla_{\theta_i} (p (\hat{L}_i(\theta_i) + \frac{\lambda}{2} \|\theta_i\|^2)) = p (\nabla \hat{L}_i(\theta_i) + \lambda \theta_i)$ if $i \in U$ and zero otherwise. Summing the contributions gives the stated expression. $\square$

Lemma 14 (Local Lipschitz of $a_i(\Theta)$). *Let the standard regularity and boundedness assumptions for the idealized MSGD dynamics hold. Then $a_i(\Theta)$ is locally Lipschitz in Θ .*

Proof. From Lemma 8, we know that the loss function ℓ is locally Lipschitz. In other words: for every compact set $\mathcal{K} \subset \mathbb{R}^{k \times d}$ there is a constant $L_{\mathcal{K}} < \infty$ such that

$$|\ell(x, \theta) - \ell(x, \theta')| \leq L_{\mathcal{K}} \|\theta - \theta'\| \quad \forall x \in B(0, R), \theta, \theta' \in \mathcal{K}.$$

Fix any compact neighborhood $\mathcal{K}$ containing both Θ and Θ' . Let

$$p_{\max} = \sup_{x \in B(0, R)} p(x), \quad L_{\mathcal{K}} \text{ as above.}$$

We will show $|a_i(\Theta) - a_i(\Theta')| \leq C_{\mathcal{K}} \|\Theta - \Theta'\|$ for some $C_{\mathcal{K}}$.

Case $m = 2$. With services $i = 1, 2$,

$$a_1(\Theta) - a_1(\Theta') = \int_{X_1(\Theta) \setminus X_1(\Theta')} p(x) dx - \int_{X_1(\Theta') \setminus X_1(\Theta)} p(x) dx,$$

so

$$|a_1(\Theta) - a_1(\Theta')| \leq p_{\max} \left[\lambda(X_1(\Theta) \setminus X_1(\Theta')) + \lambda(X_1(\Theta') \setminus X_1(\Theta)) \right].$$

For any $x \in X_1(\Theta') \setminus X_1(\Theta)$ we have $\ell(x, \theta'_1) < \ell(x, \theta'_2)$ and $\ell(x, \theta_2) < \ell(x, \theta_1)$. Hence

$$0 < \ell(x, \theta_1) - \ell(x, \theta_2) = [\ell(x, \theta_1) - \ell(x, \theta'_1)] + [\ell(x, \theta'_1) - \ell(x, \theta'_2)] + [\ell(x, \theta'_2) - \ell(x, \theta_2)].$$

Since $x \in X_1(\Theta')$, the middle term satisfies $\ell(x, \theta'_1) - \ell(x, \theta'_2) \leq 0$. The first and third terms are each bounded by $L_{\mathcal{K}} \|\Theta - \Theta'\|$ by Lipschitzness. Thus $\ell(x, \theta_1) - \ell(x, \theta_2) \leq 2L_{\mathcal{K}} \|\Theta - \Theta'\|$. Define

$$S = \{x : |\ell(x, \theta_1) - \ell(x, \theta_2)| \leq 2L_{\mathcal{K}} \|\Theta - \Theta'\|\}.$$

Since $\lambda(S) \leq (2L_{\mathcal{K}}/C) \|\Theta - \Theta'\|$ for some constant C from Assumption 3 and $X_1(\Theta') \setminus X_1(\Theta) \subset S$, we get $\lambda(X_1(\Theta') \setminus X_1(\Theta)) \leq C' \|\Theta - \Theta'\|$. The same argument applies to the other set difference, so altogether

$$|a_1(\Theta) - a_1(\Theta')| \leq 4p_{\max} L_{\mathcal{K}} \|\Theta - \Theta'\|.$$

General m . The same pairwise argument shows for each i and $j \neq i$, $\lambda(X_i(\Theta) \cap X_j(\Theta')) \leq C_{\mathcal{K}} \|\Theta - \Theta'\|$. Summing over all $j \neq i$ gives $\lambda(X_i(\Theta) \triangle X_i(\Theta')) \leq C''_{\mathcal{K}} \|\Theta - \Theta'\|$, and hence $|a_i(\Theta) - a_i(\Theta')| \leq p_{\max} C''_{\mathcal{K}} \|\Theta - \Theta'\|$.

Since all constants depend only on the compact set $\mathcal{K}$, this proves that $a_i(\Theta)$ is Lipschitz on $\mathcal{K}$, i.e. locally Lipschitz in Θ . $\square$

E Squared Loss: Performance Guarantee

Notation. Let $\{(x_i^q, y_i^q)\}_{q=1}^n$ be the probing sample with true labels y_i^q and pseudo-labels $\tilde{y}_i^q$. Here, the true labels y_i^q are hidden from the learner, and the pseudo-labels $\tilde{y}_i^q$ are observed. For any $\theta \in \mathbb{R}^d$, define

$$\widehat{L}_i(\theta) := \frac{1}{n} \sum_{q=1}^n (\langle x_i^q, \theta \rangle - \tilde{y}_i^q)^2,$$

and

$$\widehat{L}_{i,\text{true}}(\theta) := \frac{1}{n} \sum_{q=1}^n (\langle x_i^q, \theta \rangle - y_i^q)^2.$$

Additionally, define the empirical pseudo-true discrepancy

$$\Delta_i^2 := \frac{1}{n} \sum_{q=1}^n (\tilde{y}_i^q - y_i^q)^2.$$

E.1 Proof of Lemma 3

We restate the lemma for ease of reference, and then provide the proof.

Lemma 3. *For each scenario in Definitions 2–3, the probing rule $T_i(x)$ in Table 1 satisfies Assumption 5 with the stated accuracy parameter B .*

Proof. We treat each scenario in turn.

(i) Majority-good. Fix any x with $\|x\| \leq R$. Write the peer deviations $u_j := \langle x, \theta_j^0 \rangle - \langle x, \theta^* \rangle$. If strictly more than half of the peers satisfy $\|\theta_j^0 - \theta^*\| \leq r$, then at least half of the $\{u_j\}$ lie in $[-Rr, Rr]$, so the median obeys $|\tilde{y}(x) - \langle x, \theta^* \rangle| \leq Rr$ and hence

$$(\tilde{y} - y)^2 = (\tilde{y} - \langle x, \theta^* \rangle + \langle x, \theta^* \rangle - y)^2 \leq 2(\tilde{y} - \langle x, \theta^* \rangle)^2 + 2(\langle x, \theta^* \rangle - y)^2 \leq 2R^2r^2 + 2(\langle x, \theta^* \rangle - y)^2.$$

Taking expectation over $(x, y) \sim \mathcal{P}$ yields $\mathbb{E}[(\tilde{y} - y)^2] \leq 2R^2r^2 + 2\mathbb{E}[(\langle x, \theta^* \rangle - y)^2] = 2R^2r^2 + 2\epsilon$.

(ii) Market-leader. Here $\tilde{y}(x) = x^\top \theta_{j^*}$ and by assumption $\mathbb{E}[(\tilde{y} - y)^2] = \mathbb{E}[(x^\top \theta_{j^*} - y)^2] \leq \xi$, which directly verifies Accurate Probing with $B = \xi$.

(iii) Partial knowledge. Fix any x with $\|x\| \leq R$. The probing rule $T_i(x) = G$ uses median aggregation over the subset $G \subseteq [m] \setminus \{i\}$. Write the peer deviations for $j \in G$: $u_j := \langle x, \theta_j^0 \rangle - \langle x, \theta^* \rangle$. Since all learners in G satisfy $\|\theta_j^0 - \theta^*\| \leq r$, all deviations $\{u_j\}_{j \in G}$ lie in $[-Rr, Rr]$. Because $|G| > (m-1)/2$, the set G contains more than half of the peers (excluding i), and thus the median over G obeys $|\tilde{y}(x) - \langle x, \theta^* \rangle| \leq Rr$. The remainder of the proof follows identically to case (i):

$$(\tilde{y} - y)^2 = (\tilde{y} - \langle x, \theta^* \rangle + \langle x, \theta^* \rangle - y)^2 \leq 2(\tilde{y} - \langle x, \theta^* \rangle)^2 + 2(\langle x, \theta^* \rangle - y)^2 \leq 2R^2r^2 + 2(\langle x, \theta^* \rangle - y)^2.$$

Taking expectation over $(x, y) \sim \mathcal{P}$ yields $\mathbb{E}[(\tilde{y} - y)^2] \leq 2R^2r^2 + 2\mathbb{E}[(\langle x, \theta^* \rangle - y)^2] = 2R^2r^2 + 2\epsilon$.

(iv) Preference-aware. By the probing rule $T_i(x) = \{\pi(x)\}$, we have $\tilde{y}(x) = x^\top \bar{\theta}_{\pi(x)}$, so

$$\mathbb{E}[(\tilde{y} - y)^2] = \sum_{i=1}^m \alpha_i \mathbb{E}_{(x,y) \sim \mathcal{P}_i} [(y - x^\top \bar{\theta}_i)^2].$$

By optimality of $\bar{\theta}_i$ for the ERM objective on $\mathcal{P}_i$, for every i and any θ (in particular θ^*),

$$\mathbb{E}_{\mathcal{P}_i} [(y - x^\top \bar{\theta}_i)^2] \leq \mathbb{E}_{\mathcal{P}_i} [(y - x^\top \theta^*)^2].$$

Summing over i with weights α_i gives

$$\mathbb{E}[(\tilde{y} - y)^2] \leq \sum_{i=1}^m \alpha_i \mathbb{E}_{\mathcal{P}_i} [(y - x^\top \theta^*)^2] = \epsilon.$$

□

E.2 Proof of Theorem 4

Proof of Corollary 1. Set $\lambda = \epsilon/\|\theta^\star\|^2$ in Lemma 15. Define

$$S(\kappa, \lambda) := (4b_\star + 6C_0^2)\sqrt{2\log(2/\kappa)} + 4(Y_{\max} + B_\theta R)B_\theta R + b_u\sqrt{2\log(2/\kappa)},$$

where B_θ and b_u are the λ -dependent quantities from Lemma 15. If

$$n \geq \underline{n} := \frac{S(\kappa, \lambda)^2}{\epsilon^2},$$

then the concentration terms in Lemma 15 satisfy $S(\kappa, \lambda)/\sqrt{n} \leq \epsilon$. With this choice and $\lambda\|\theta^\star\|^2 = \epsilon$, the explicit bound yields

$$\mathcal{R}(\tilde{\theta}_i) \leq 6B + \left(4 + \frac{2}{p}\right)\epsilon + \epsilon + \epsilon = 6B + \left(6 + \frac{2}{p}\right)\epsilon,$$

which is $O\left(\left(\frac{p+1}{p}\right)\epsilon + B\right)$.

To see the stated scaling of $\underline{n}$, note that $B_\theta = \Theta(1/\sqrt{\lambda})$ and $b_u = (Y_{\max} + B_\theta R)^2 = O(1/\lambda)$ with constants depending on $(R, Y_{\max}, M_0, p)$. To surface the dominant dependence on $(R, Y_{\max}, M_0, p)$, write $A := Y_{\max}^2 + pR^2M_0^2$ so that

$$B_\theta R = R\sqrt{\frac{2A}{\lambda p}} \quad \text{and} \quad (B_\theta R)^2 = \frac{2R^2A}{\lambda p}.$$

The leading terms (in terms of ϵ) in $S(\kappa, \lambda)$ scale as $(B_\theta R)^2$, so one can take

$$\underline{n} = O\left(\frac{R^4\|\theta^\star\|^4}{\epsilon^4} \left(\frac{Y_{\max}^2}{p} + R^2M_0^2\right)^2 \log \frac{1}{\kappa}\right),$$

where we suppress lower-order terms in ϵ coming from $b_\star$ and C_0 . $\square$

Theorem 4. *Let Assumptions 1–5 hold. For any $\kappa \in (0, 1)$, with probability at least $1 - \kappa$, we have for probing learner $i \in [m]$, for the squared loss:*

$$\mathcal{R}(\tilde{\theta}_i) \leq O\left(\left(\frac{p+1}{p}\right)\epsilon + B + \lambda\|\theta^\star\|^2 + \frac{(p+1)C_{\text{gen}}}{p\lambda}\sqrt{\frac{\log(1/\kappa)}{n}}\right).$$

where $C_{\text{gen}}(R, Y_{\max}, \|\theta^\star\|, \max_{j \neq i} \|\theta_j^0\|)$ is stated explicitly in the Appendix.

Proof. By Lemma 15, for any $\kappa \in (0, 1)$ and with probability at least $1 - \kappa$,

$$\mathcal{R}(\tilde{\theta}_i) \leq 6B + \left(4 + \frac{2}{p}\right)\epsilon + \lambda\|\theta^\star\|^2 + \frac{T(\kappa)}{\sqrt{n}},$$

where every $n^{-1/2}$ contribution has been grouped into

$$T(\kappa) := (4b_\star + 6C_0^2)\sqrt{2\log(2/\kappa)} + b_u\sqrt{2\log(2/\kappa)} + 4(Y_{\max} + B_\theta R)B_\theta R. \quad (21)$$

The terms that depend on λ and p are through $B_\theta = \sqrt{2(Y_{\max}^2 + pY_{\max, \text{probe}}^2)/(\lambda p)}$, which also appears inside $b_u = (Y_{\max} + B_\theta R)^2$. Expanding the B_θ -dependent pieces, from (21),

$$T(\kappa) = (4b_\star + 6C_0^2)\sqrt{2\log(2/\kappa)} + \underbrace{\left[(Y_{\max} + B_\theta R)^2\sqrt{2\log(2/\kappa)} + 4(Y_{\max} + B_\theta R)B_\theta R\right]}_{B_\theta\text{-dependent}}.$$

Thus the λ -independent part is already $O(\sqrt{\log(1/\kappa)})$, so it remains to control the bracketed term. Expanding gives

$$(Y_{\max} + B_{\theta}R)^2 \sqrt{2 \log(2/\kappa)} + 4(Y_{\max} + B_{\theta}R)B_{\theta}R = O\left(\sqrt{\log(1/\kappa)} [Y_{\max}^2 + Y_{\max} B_{\theta}R + B_{\theta}^2 R^2]\right).$$

Plugging in the value of B_{θ} from Lemma 20 and assuming $\lambda < 1$, we can combine with the (λ, p) -independent part of $T(\kappa)$ to get

$$C_{\text{gen}} = R(Y_{\max}^2 + Y_{\max, \text{probe}}^2) + 4b_{\star} + 6C_0^2$$

□

Lemma 15 (Squared-loss performance Full Statement). *Let Assumption 5 hold with parameter B . Then, for any $\kappa \in (0, 1)$, with probability at least $1 - \kappa$ over the probing sample, every stationary point $\tilde{\Theta}$ of MSGD-P satisfies, for probing learner $i \in [m]$,*

$$\begin{aligned} \mathcal{R}(\tilde{\theta}_i) \leq & 6B + \left(4 + \frac{2}{p}\right) \epsilon + \lambda \|\theta^{\star}\|^2 \\ & + \frac{(4b_{\star} + 6C_0^2) \sqrt{2 \log(2/\kappa)}}{\sqrt{n}} + \frac{4(Y_{\max} + B_{\theta}R) B_{\theta}R}{\sqrt{n}} + b_u \sqrt{\frac{2 \log(2/\kappa)}{n}}, \end{aligned}$$

where the constants are defined as follows:

- $C_0 := Y_{\max} + Y_{\max, \text{probe}}$.
- $B_{\theta} := \sqrt{\frac{2(Y_{\max}^2 + p Y_{\max, \text{probe}}^2)}{\lambda p}}$ is a radius (independent of $\theta^{\star}$) with $Y_{\max, \text{probe}} := R M_0$, $M_0 := \max_{j \neq i} \|\theta_j^0\|$, which bounds the pseudo-label magnitudes via $|\tilde{y}_i^q| \leq Y_{\max, \text{probe}}$.
- $b_{\star} := (Y_{\max} + R \|\theta^{\star}\|)^2$.
- $b_u := (Y_{\max} + B_{\theta}R)^2$.

Proof. We work on the event $\mathcal{E}_u \cap \mathcal{E}_{\star}$ where $\mathcal{E}_u$ is the event of Lemma 18 (probability $\geq 1 - \kappa/2$ with B_{θ} from Lemma 20) and $\mathcal{E}_{\star}$ is the event of Lemma 17 (probability $\geq 1 - \kappa/2$). By a union bound,

$$\mathbb{P}(\mathcal{E}_u \cap \mathcal{E}_{\star}) \geq 1 - \kappa.$$

Bound on $\hat{L}_i(\tilde{\theta}_i)$. By stationarity, $\tilde{\theta}_i$ minimizes

$$\hat{\Phi}_i^{\text{probe}}(\theta; \tilde{\Theta}) = \tau \alpha_i \mathbb{E}_{\mathcal{P}_i}[\ell(z, \theta)] + (1 - \tau) a_i(\tilde{\Theta}) \mathbb{E}_{\mathcal{D}_i(\tilde{\Theta})}[\ell(z, \theta)] + p \hat{L}_i(\theta) + \frac{\lambda p}{2} \|\theta\|^2.$$

Thus, comparing the objective at $\tilde{\theta}_i$ and $\theta^{\star}$,

$$p \hat{L}_i(\tilde{\theta}_i) \leq \tau \alpha_i \mathbb{E}_{\mathcal{P}_i}[\ell(z, \theta^{\star})] + (1 - \tau) a_i(\tilde{\Theta}) \mathbb{E}_{\mathcal{D}_i(\tilde{\Theta})}[\ell(z, \theta^{\star})] + p \hat{L}_i(\theta^{\star}) + \frac{\lambda p}{2} \|\theta^{\star}\|^2.$$

Since

$$\tau \alpha_i \mathbb{E}_{\mathcal{P}_i} \ell(\theta^{\star}) + (1 - \tau) a_i(\tilde{\Theta}) \mathbb{E}_{\mathcal{D}_i(\tilde{\Theta})} \ell(\theta^{\star}) \leq \epsilon,$$

dividing by p , we obtain

$$\widehat{L}_i(\tilde{\theta}_i) \leq \frac{\epsilon}{p} + \widehat{L}_i(\theta^*) + \frac{\lambda}{2} \|\theta^*\|^2.$$

By Lemma 16 and Lemma 17,

$$\widehat{L}_i(\theta^*) \leq 2 \widehat{L}_{i,\text{true}}(\theta^*) + 2 \Delta_i^2 \leq 2 \epsilon + 2 \Lambda_\star + 2 \Delta_i^2.$$

Hence

$$\widehat{L}_i(\tilde{\theta}_i) \leq \frac{\epsilon}{p} + 2 \epsilon + 2 \Lambda_\star + 2 \Delta_i^2 + \frac{\lambda}{2} \|\theta^*\|^2.$$

Bound on $\widehat{L}_{i,\text{true}}(\tilde{\theta}_i)$. Applying Lemma 16 again gives

$$\widehat{L}_{i,\text{true}}(\tilde{\theta}_i) \leq 2 \widehat{L}_i(\tilde{\theta}_i) + 2 \Delta_i^2.$$

Substituting the bound from Step 1,

$$\widehat{L}_{i,\text{true}}(\tilde{\theta}_i) \leq \left(\frac{2}{p} + 4\right) \epsilon + 4 \Lambda_\star + 6 \Delta_i^2 + \lambda \|\theta^*\|^2.$$

By Lemma 19, on $\mathcal{E}_\star$,

$$\Delta_i^2 \leq B + C_0^2 \sqrt{\frac{2 \log(2/\kappa)}{n}}.$$

Therefore,

$$\widehat{L}_{i,\text{true}}(\tilde{\theta}_i) \leq 6B + \left(\frac{2}{p} + 4\right) \epsilon + 4 \Lambda_\star + 6C_0^2 \sqrt{\frac{2 \log(2/\kappa)}{n}} + \lambda \|\theta^*\|^2.$$

Bound on $\mathcal{R}(\tilde{\theta}_i)$. Finally, on $\mathcal{E}_u$ (with $\|\tilde{\theta}_i\| \leq B_\theta$ from Lemma 20),

$$\mathcal{R}(\tilde{\theta}_i) \leq \widehat{L}_{i,\text{true}}(\tilde{\theta}_i) + \Lambda_u.$$

□

Alternative scaling with explicit λ . If we keep λ explicit and only choose n to control the concentration terms in Lemma 15, then for any target tolerance $\delta > 0$ it suffices to take

$$n \geq \bar{n}(\lambda, \delta) := \frac{S(\kappa, \lambda)^2}{\delta^2},$$

which yields

$$\mathcal{R}(\tilde{\theta}_i) \leq 6B + \left(4 + \frac{2}{p}\right) \epsilon + \lambda \|\theta^*\|^2 + \delta.$$

Writing $Y_{\text{max,probe}} = RM_0$ and $A := Y_{\text{max}}^2 + pR^2M_0^2$, we have

$$B_\theta = \sqrt{\frac{2A}{\lambda p}} \quad \text{and} \quad (Y_{\text{max}} + B_\theta R)^2 \leq 2Y_{\text{max}}^2 + \frac{4R^2A}{\lambda p}.$$

Using $b_\star = (Y_{\text{max}} + R\|\theta^*\|)^2$ and $C_0 = Y_{\text{max}} + RM_0$, this implies

$$S(\kappa, \lambda)^2 = O\left(\log \frac{1}{\kappa}\right) \left[(Y_{\text{max}} + R\|\theta^*\|)^4 + (Y_{\text{max}} + RM_0)^4 + \frac{R^4 A^2}{\lambda^2 p^2} \right],$$

and therefore a sufficient choice is

$$n = O\left(\frac{\log(1/\kappa)}{\delta^2} \left[(Y_{\max} + R\|\theta^*\|)^4 + (Y_{\max} + RM_0)^4 + \frac{R^4(Y_{\max}^2 + pR^2M_0^2)^2}{\lambda^2 p^2} \right] \right).$$

Setting $\delta = \epsilon$ yields the same bias expression as above, but now with n explicit in $(R, Y_{\max}, M_0, p, \lambda, \epsilon, \kappa)$. $\square$

Lemma 16. *We have*

$$\widehat{L}_{i,\text{true}}(\theta) \leq \left(\sqrt{\widehat{L}_i(\theta)} + \Delta_i \right)^2 \quad \text{and} \quad \widehat{L}_i(\theta) \leq \left(\sqrt{\widehat{L}_{i,\text{true}}(\theta)} + \Delta_i \right)^2.$$

In particular, $\widehat{L}_{i,\text{true}}(\theta) \leq 2\widehat{L}_i(\theta) + 2\Delta_i^2$ and $\widehat{L}_i(\theta) \leq 2\widehat{L}_{i,\text{true}}(\theta) + 2\Delta_i^2$.

Proof of Lemma 16. Let $X \in \mathbb{R}^{n \times d}$ be the design matrix, $\mathbf{a} := X\theta - \tilde{\mathbf{y}}$ and $\mathbf{b} := \tilde{\mathbf{y}} - \mathbf{y}$. Then

$$X\theta - \mathbf{y} = \mathbf{a} + \mathbf{b}.$$

By the triangle inequality,

$$\sqrt{\widehat{L}_{i,\text{true}}(\theta)} = \frac{\|X\theta - \mathbf{y}\|}{\sqrt{n}} \leq \frac{\|\mathbf{a}\|}{\sqrt{n}} + \frac{\|\mathbf{b}\|}{\sqrt{n}} = \sqrt{\widehat{L}_i(\theta)} + \Delta_i.$$

Squaring yields the first inequality; the second is analogous. $\square$

Lemma 17 (One-point deviation at θ^*). *For any $\kappa \in (0, 1)$, with probability at least $1 - \kappa/2$,*

$$|\widehat{L}_{i,\text{true}}(\theta^*) - \mathcal{R}(\theta^*)| \leq \Lambda_\star \quad \text{where} \quad \Lambda_\star := b_\star \sqrt{\frac{2\log(2/\kappa)}{n}}, \quad b_\star := (Y_{\max} + R\|\theta^*\|)^2.$$

Proof of Lemma 17. For each q , define the random variable

$$Z_q := (\langle \tilde{x}_i^q, \theta^* \rangle - y_i^q)^2.$$

Since $|y| \leq Y_{\max}$ and $\|x\| \leq R$ a.s., we have

$$Z_q \in [0, (Y_{\max} + R\|\theta^*\|)^2] = [0, b_\star].$$

By Hoeffding's inequality, with probability at least $1 - \kappa/2$,

$$\left| \frac{1}{n} \sum_{q=1}^n Z_q - \mathbb{E}[Z] \right| \leq b_\star \sqrt{\frac{2\log(2/\kappa)}{n}}.$$

$\square$

Lemma 18 (Uniform deviation over $\{\|\theta\| \leq B_\theta\}$). *Fix any radius $B_\theta > 0$ and let $\kappa \in (0, 1)$. Define the squared-loss class*

$$\mathcal{F}_{B_\theta} := \left\{ f_\theta(x, y) = (y - \langle x, \theta \rangle)^2 \mid \|\theta\| \leq B_\theta \right\}.$$

Then, with probability at least $1 - \kappa/2$ (over the draw of the probing sample of size n),

$$\sup_{\|\theta\| \leq B_\theta} |\widehat{L}_{i,\text{true}}(\theta) - \mathcal{R}(\theta)| \leq \Lambda_u,$$

where

$$\Lambda_u := \frac{4(Y_{\max} + B_\theta R) B_\theta R}{\sqrt{n}} + b_u \sqrt{\frac{2\log(2/\kappa)}{n}}, \quad b_u := (Y_{\max} + B_\theta R)^2.$$

Proof. We apply Wainwright's Theorem 4.10.

Uniform bound b for $\mathcal{F}_{B_\theta}$. For any $\|\theta\| \leq B_\theta$ and any (x, y) with $\|x\| \leq R$, $|y| \leq Y_{\max}$, we have

$$|y - \langle x, \theta \rangle| \leq |y| + |\langle x, \theta \rangle| \leq Y_{\max} + \|x\| \|\theta\| \leq Y_{\max} + B_\theta R.$$

Hence

$$0 \leq f_\theta(x, y) = (y - \langle x, \theta \rangle)^2 \leq (Y_{\max} + B_\theta R)^2 =: b_u.$$

Thus $\mathcal{F}_{B_\theta}$ is b_u -uniformly bounded.

Bound $R_n(\mathcal{F}_{B_\theta})$ via contraction to the linear class. Fix a sample $S = \{(x_i, y_i)\}_{i=1}^n$ and let $\sigma_1, \dots, \sigma_n$ be i.i.d. Rademacher signs. The empirical Rademacher average of $\mathcal{F}_{B_\theta}$ on S is

$$\text{Rad}_S(\mathcal{F}_{B_\theta}) := \mathbb{E}_\sigma \left[\sup_{\|\theta\| \leq B_\theta} \frac{1}{n} \sum_{i=1}^n \sigma_i (y_i - \langle x_i, \theta \rangle)^2 \right].$$

Define $u_{i,\theta} := \langle x_i, \theta \rangle$ and, for each i , the recentered function

$$\psi_i(u) := (y_i - u)^2 - (y_i - 0)^2 = u^2 - 2y_i u, \quad \psi_i(0) = 0.$$

Since $\mathbb{E}[\sigma_i] = 0$, the constant offsets $\{(y_i - 0)^2\}$ vanish in the Rademacher average. Thus

$$\text{Rad}_S(\mathcal{F}_{B_\theta}) = \mathbb{E}_\sigma \left[\sup_{\|\theta\| \leq B_\theta} \frac{1}{n} \sum_{i=1}^n \sigma_i \psi_i(u_{i,\theta}) \right].$$

We now bound the Lipschitz constants of ψ_i over the relevant range. For any $u, v \in [-B_\theta R, B_\theta R]$, $|\psi_i(u) - \psi_i(v)| = |(u-v)[(u-y_i) + (v-y_i)]| \leq |u-v| (|u-y_i| + |v-y_i|) \leq 2(Y_{\max} + B_\theta R) |u-v|.$

Hence each ψ_i is L -Lipschitz with

$$L := 2(Y_{\max} + B_\theta R), \quad \psi_i(0) = 0.$$

By the Ledoux–Talagrand contraction inequality (applied elementwise to $\{\psi_i\}$ with common Lipschitz constant L and $\psi_i(0) = 0$), we have

$$\text{Rad}_S(\mathcal{F}_{B_\theta}) \leq L \cdot \mathbb{E}_\sigma \left[\sup_{\|\theta\| \leq B_\theta} \frac{1}{n} \sum_{i=1}^n \sigma_i u_{i,\theta} \right] = L \cdot \text{Rad}_S(\mathcal{G}_{B_\theta}),$$

where $\mathcal{G}_{B_\theta} := \{(x, y) \mapsto \langle x, \theta \rangle : \|\theta\| \leq B_\theta\}$ is the linear class (note: dependence on y disappears).

The empirical Rademacher average of $\mathcal{G}_{B_\theta}$ on S is standard:

$$\text{Rad}_S(\mathcal{G}_{B_\theta}) = \mathbb{E}_\sigma \left[\sup_{\|\theta\| \leq B_\theta} \frac{1}{n} \sum_{i=1}^n \sigma_i \langle x_i, \theta \rangle \right] = \frac{B_\theta}{n} \mathbb{E}_\sigma \left\| \sum_{i=1}^n \sigma_i x_i \right\| \leq \frac{B_\theta}{n} \sqrt{\mathbb{E}_\sigma \left\| \sum_{i=1}^n \sigma_i x_i \right\|^2} = \frac{B_\theta}{n} \sqrt{\sum_{i=1}^n \|x_i\|^2}.$$

Using $\|x_i\| \leq R$, we conclude $\text{Rad}_S(\mathcal{G}_{B_\theta}) \leq B_\theta R / \sqrt{n}$. Combining,

$$\text{Rad}_S(\mathcal{F}_{B_\theta}) \leq L \cdot \text{Rad}_S(\mathcal{G}_{B_\theta}) \leq 2(Y_{\max} + B_\theta R) \frac{B_\theta R}{\sqrt{n}}.$$

Taking expectation over the sample (or simply noting that this bound holds for any S satisfying $\|x_i\| \leq R$) yields

$$R_n(\mathcal{F}_{B_\theta}) \leq 2(Y_{\max} + B_\theta R) \frac{B_\theta R}{\sqrt{n}}.$$

Apply Theorem 4.10 and choose δ . By Theorem 4.10, for any $\delta > 0$,

$$\sup_{\|\theta\| \leq B_\theta} |\widehat{L}_{i,\text{true}}(\theta) - \mathcal{R}(\theta)| \leq 2 R_n(\mathcal{F}_{B_\theta}) + \delta \leq \frac{4(Y_{\max} + B_\theta R) B_\theta R}{\sqrt{n}} + \delta,$$

with probability at least $1 - \exp(-\frac{n\delta^2}{2b_u^2})$. Set $\delta = b_u \sqrt{\frac{2\log(2/\kappa)}{n}}$ so that $\exp(-\frac{n\delta^2}{2b_u^2}) = \exp(-\log(2/\kappa)) = \kappa/2$. The stated deviation bound Λ_u follows. $\square$

Lemma 19 (Concentration of Δ_n^2 under Accurate Probing). *Assume that $|\tilde{y}| \leq Y_{\max, \text{probe}}$ almost surely (e.g., $Y_{\max, \text{probe}} = R \max_j \|\theta_j^0\|$ for linear peers at initialization). Then for any $\kappa \in (0, 1)$, with probability at least $1 - \kappa/2$,*

$$\Delta_n^2 \leq B + C_0^2 \sqrt{\frac{2\log(2/\kappa)}{n}}, \quad \text{where } C_0 := Y_{\max} + Y_{\max, \text{probe}}.$$

Proof. Let $W = (\tilde{y} - y)^2$. By the boundedness assumptions, $0 \leq W \leq (Y_{\max} + Y_{\max, \text{probe}})^2 =: U$ almost surely, and $\mathbb{E}[W] \leq B$ by Accurate Probing. By Hoeffding's inequality,

$$\Pr\left(\frac{1}{n} \sum_{q=1}^n W_q - \mathbb{E}[W] \geq t\right) \leq \exp\left(-\frac{2nt^2}{U^2}\right).$$

Choosing $t = U \sqrt{\frac{2\log(2/\kappa)}{n}}$ yields $\Pr(\Delta_n^2 \geq \mathbb{E}[W] + U \sqrt{\frac{2\log(2/\kappa)}{n}}) \leq \kappa/2$. Using $\mathbb{E}[W] \leq B$ gives the stated bound with $C_0^2 = U$. $\square$

Lemma 20 (A priori norm bound and choice of B_θ). *For any probing learner $i \in U$,*

$$\|\tilde{\theta}_i\| \leq B_\theta \quad \text{with} \quad B_\theta := \sqrt{\frac{2(Y_{\max}^2 + p Y_{\max, \text{probe}}^2)}{\lambda p}}.$$

Proof of Lemma 20. At any stationary point $\tilde{\Theta}$, $\tilde{\theta}_i$ minimizes the convex function

$$\hat{\Phi}_i^{\text{probe}}(\theta; \tilde{\Theta}) := \tau \alpha_i \mathbb{E}_{z \sim \mathcal{P}_i}[\ell(z, \theta)] + (1 - \tau) a_i(\tilde{\Theta}) \mathbb{E}_{z \sim \mathcal{D}_i(\tilde{\Theta})}[\ell(z, \theta)] + p \widehat{L}_i(\theta) + \frac{\lambda p}{2} \|\theta\|^2.$$

By optimality of $\tilde{\theta}_i$ for $\hat{\Phi}_i^{\text{probe}}(\cdot; \tilde{\Theta})$,

$$\hat{\Phi}_i^{\text{probe}}(\tilde{\theta}_i; \tilde{\Theta}) \leq \hat{\Phi}_i^{\text{probe}}(0; \tilde{\Theta}).$$

Using squared loss and $|y| \leq Y_{\max}$, we have

$$\widehat{L}_i(0) = \frac{1}{n} \sum_{q=1}^n (\tilde{y}_i^q)^2 \leq Y_{\max, \text{probe}}^2.$$

Hence,

$$\frac{\lambda p}{2} \|\tilde{\theta}_i\|^2 \leq \tau \alpha_i Y_{\max}^2 + (1 - \tau) a_i(\tilde{\Theta}) Y_{\max}^2 + p \widehat{L}_i(0) \leq Y_{\max}^2 + p Y_{\max, \text{probe}}^2.$$

Thus

$$\|\tilde{\theta}_i\| \leq \sqrt{\frac{2(Y_{\max}^2 + p Y_{\max, \text{probe}}^2)}{\lambda p}} = B_\theta. \quad \square$$

F Cross-Entropy Loss: Performance Guarantee

Notation. In this section, we provide a high-probability upper bound on the full-population cross-entropy risk for a probing learner in MSGD-P at a stationary point. The proof mirrors the squared-loss analysis, with key differences arising from the geometry of the softmax and cross-entropy.

Throughout this section, let K denote the number of classes. Let $\theta \equiv W \in \mathbb{R}^{K \times d}$ be the matrix of class-wise linear weights and define logits $z(x) = Wx \in \mathbb{R}^K$ and predicted probabilities $q_W(x) = \text{softmax}(z(x))$. The multiclass cross-entropy loss is

$$\ell_{\text{CE}}((x, y), W) = - \sum_{c=1}^K y_c \log q_W(x)_c = -\log q_W(x)_Y,$$

where $y \in \{e_1, \dots, e_K\}$ is one-hot and Y is the true class index. We assume $\|x\| \leq R$ almost surely.

For a probing learner $i \in U$ (with probing weight $p > 0$), we define its augmented objective at Θ (as per MSGD-P) by

$$\begin{aligned} \Phi_i(W; \Theta) &= \tau \alpha_i \mathbb{E}_{(x, y) \sim \mathcal{P}_i} [\ell_{\text{CE}}((x, y), W)] \\ &\quad + (1 - \tau) a_i(\Theta) \mathbb{E}_{(x, y) \sim \mathcal{D}_i(\Theta)} [\ell_{\text{CE}}((x, y), W)] \\ &\quad + p \hat{L}_{\text{probe}}(W) + \frac{\lambda p}{2} \|W\|_F^2, \end{aligned}$$

where

$$\hat{L}_{\text{probe}}(W) = \frac{1}{n} \sum_{q=1}^n \ell_{\text{CE}}((\tilde{x}_i^q, \tilde{y}_i^q), W) = -\frac{1}{n} \sum_{q=1}^n \sum_{c=1}^K \tilde{y}_{i,c}^q \log q_W(\tilde{x}_i^q)_c.$$

Here $\{(\tilde{x}_i^q)\}_{q=1}^n$ are probing covariates drawn i.i.d. from $\mathcal{P}_X$, and $\{\tilde{y}_i^q\}_{q=1}^n$ are soft pseudo-labels. For analysis, let y_i^q be the (hidden) true labels of $\tilde{x}_i^q$, and define the empirical “true” probing CE:

$$\hat{L}_{\text{true}}(W) = \frac{1}{n} \sum_{q=1}^n \ell_{\text{CE}}((\tilde{x}_i^q, y_i^q), W) = -\frac{1}{n} \sum_{q=1}^n \log q_W(\tilde{x}_i^q)_{y_i^q}.$$

Let $\theta^* \in \arg \min_W \mathcal{R}(W)$ be a (finite-norm) population minimizer of the unregularized CE risk $\mathcal{R}(W) = \mathbb{E}[\ell_{\text{CE}}((x, y), W)]$, with $\epsilon := \mathcal{R}(\theta^*)$ and $\|\theta^*\|_F =: M_\star < \infty$.

Aggregation Rule. we define peer logits $z_j(x, \Theta) = \theta_j^0 x \in \mathbb{R}^K$ and the coordinatewise median of logits

$$\tilde{z}_c(x, \Theta) := \text{median}(z_{j,c}(x, \Theta) : j \in T_i(x)), \quad c \in [K].$$

Then, we set the probing label as $y_{\text{agg},i}(x, \Theta) = \text{softmax}(\tilde{z}(x, \Theta))$.

F.1 Accurate Probing Assumption and Proofs

Assumption 6 (Accurate Probing for CE). *There exists $B_{\text{CE}} \geq 0$ such that*

$$\mathbb{E}[\text{CE}(y, \tilde{y})] = \mathbb{E}[-\log \tilde{y}_Y] \leq B_{\text{CE}},$$

where the expectation is over $(x, y) \sim \mathcal{P}$ and $\tilde{y} = \tilde{y}(x)$ produced by the robust aggregator.

We now present sufficient conditions ensuring Assumption 6.

Scenario	$T_i(x)$	B
Majority-good	$[m]$	$\epsilon + 2Rr$
Market-leader	$\{j^*\}$	ξ
Partial knowledge	G	$\epsilon + 2Rr$
Preference-aware	$\{\pi(x)\}$	ϵ

Table 2: Probing accuracy parameters for cross-entropy loss.

Lemma 21 (Sufficient conditions for Accurate Probing (CE)). *Assumption 6 holds in each of the following scenarios: (i) **Majority-good**. Suppose strictly more than half of peers satisfy $\|\theta_j^0 - \theta^*\|_F \leq r$. Then*

$$B_{\text{CE}} = \mathbb{E}[-\log \tilde{y}_Y] \leq \epsilon + 2Rr.$$

*(ii) **Market-leader**. Suppose learner i probes a single peer j^* with $\mathbb{E}[\ell_{\text{CE}}((x, y), \theta_{j^*}^0)] \leq \xi$. If $\tilde{y}(x) = q_{\theta_{j^*}^0}(x)$, then $B_{\text{CE}} \leq \xi$.*

*(iii) **Partial knowledge**. If a known subset G of peers satisfies $|G| > (m-1)/2$ and $\|\theta_j^0 - \theta^*\|_F \leq r$ for all $j \in G$, and the aggregator uses the coordinatewise median over G , the same bound as in case (i) holds.*

*(iv) **Preference-aware**. Suppose each peer j solves the ERM on its preference partition $\mathcal{P}_j$:*

$$\bar{\theta}_j \in \arg \min_W \mathbb{E}_{(x,y) \sim \mathcal{P}_j} [\ell_{\text{CE}}((x, y), W)].$$

If learner i probes $\tilde{y}(x) = q_{\bar{\theta}_{\pi(x)}}(x)$, then

$$B_{\text{CE}} \leq \sum_{j=1}^K \alpha_j \mathbb{E}_{\mathcal{P}_j} [\ell_{\text{CE}}((x, y), \bar{\theta}_j)] \leq \epsilon.$$

Proof. We treat each scenario in turn.

(i) Majority-good. Fix any x with $\|x\| \leq R$ and write $z^*(x) = \theta^* x$. For any good peer j (i.e., $\|\theta_j^0 - \theta^*\|_F \leq r$) and any class c ,

$$|z_{j,c}(x) - z_c^*(x)| = |(\theta_{j,c}^0 - \theta_c^*)^\top x| \leq \|\theta_{j,c}^0 - \theta_c^*\|_2 \cdot \|x\| \leq \|\theta_j^0 - \theta^*\|_F \cdot \|x\| \leq rR.$$

Since strictly more than half of the $\{z_{j,c}(x)\}_{j \neq i}$ lie in $[z_c^*(x) - rR, z_c^*(x) + rR]$, their coordinatewise median must also lie in this interval. Hence

$$\|\tilde{z}(x) - z^*(x)\|_\infty \leq rR.$$

Let $f(z) = -\log \text{softmax}(z)_Y$. By Lemma 22,

$$-\log \tilde{y}_Y = f(\tilde{z}(x)) \leq f(z^*(x)) + 2\|\tilde{z}(x) - z^*(x)\|_\infty \leq -\log q_Y^*(x) + 2rR,$$

where $q^* = \text{softmax}(z^*)$. Taking expectations over $(x, y) \sim \mathcal{P}$, we get

$$B_{\text{CE}} = \mathbb{E}[-\log \tilde{y}_Y] \leq \mathbb{E}[-\log q_Y^*(x)] + 2rR = \epsilon + 2rR.$$

(ii) Market-leader. Direct: $B_{\text{CE}} = \mathbb{E}[-\log \tilde{y}_Y] = \mathbb{E}[\ell_{\text{CE}}((x, y), \theta_{j^*}^0)] \leq \xi$.

(iii) **Partial knowledge.** The proof is identical to case (i), restricting the median to G . Since $|G| > (m-1)/2$ and all learners in G satisfy $\|\theta_j^0 - \theta^*\|_F \leq r$, the coordinatewise median over G satisfies the same bound.

(iv) **Preference-aware.** By optimality of $\bar{\theta}_j$ for the ERM objective,

$$\mathbb{E}_{\mathcal{P}_j} [\ell_{\text{CE}}((x, y), \bar{\theta}_j)] \leq \mathbb{E}_{\mathcal{P}_j} [\ell_{\text{CE}}((x, y), \theta^*)].$$

Summing over j with weights α_j gives

$$B_{\text{CE}} \leq \sum_j \alpha_j \mathbb{E}_{\mathcal{P}_j} \ell_{\text{CE}}(\bar{\theta}_j) \leq \sum_j \alpha_j \mathbb{E}_{\mathcal{P}_j} \ell_{\text{CE}}(\theta^*) = \epsilon.$$

□

Lemma 22 (CE is 2-Lipschitz in logits under $\|\cdot\|_\infty$). *Fix a class $Y \in [K]$ and define $f(z) := -\log \text{softmax}(z)_Y = -z_Y + \log \text{sumexp}(z)$ for $z \in \mathbb{R}^K$. Then for any $z, z' \in \mathbb{R}^K$,*

$$|f(z') - f(z)| \leq 2 \|z' - z\|_\infty.$$

Proof. We have $\nabla f(z) = q(z) - e_Y$, where $q(z) = \text{softmax}(z)$ and e_Y is the one-hot at Y . By the mean value inequality with dual norms $(\|\cdot\|_\infty, \|\cdot\|_1)$,

$$|f(z') - f(z)| \leq \sup_{\xi \in [z, z']} \|\nabla f(\xi)\|_1 \cdot \|z' - z\|_\infty = \sup_{\xi} \sum_{k=1}^K |q_k(\xi) - (e_Y)_k| \cdot \|z' - z\|_\infty.$$

Since $y = e_Y$ is one-hot and $q \in \Delta^{K-1}$, $\sum_{k=1}^K |q_k - (e_Y)_k| = |q_Y - 1| + \sum_{k \neq Y} q_k = (1 - q_Y) + (1 - q_Y) \leq 2$. Hence $|f(z') - f(z)| \leq 2 \|z' - z\|_\infty$. □

F.2 Performance Bound

We will use two probability floors:

$$\gamma_\star := \frac{1}{K} \exp(-2R \|\theta^*\|_F), \quad \Gamma_\star := \log \frac{1}{\gamma_\star} = \log K + 2R \|\theta^*\|_F,$$

and

$$B_\theta := \sqrt{\frac{2(1+p) \log K}{\lambda p}}, \quad \gamma_B := \frac{1}{K} \exp(-2R B_\theta), \quad \Gamma_B := \log \frac{1}{\gamma_B} = \log K + 2R B_\theta.$$

Corollary 2 (Big-O summary). *Under the assumptions of Theorem 5, with probability at least $1 - \kappa$,*

$$\mathcal{R}(\tilde{\theta}_i) \leq O\left(\left(\frac{p+1}{p}\right)\epsilon + \lambda \|\theta^*\|_F^2 + C_{\text{bias}} \sqrt{B_{\text{CE}}} + \frac{(p+1) C_{\text{gen}}}{p} \sqrt{\frac{\log(1/\kappa)}{\lambda n}}\right),$$

where the constant $C_{\text{bias}} = R \left(\|\theta^*\|_F + \sqrt{2 \log K / \lambda} \right)$ and $C_{\text{gen}} = C_{\text{gen}}(R, \|\theta^*\|, \log(K))$ is specified in the appendix.

Theorem 5 (CE performance bound with probing). *Let Assumptions 1, 2, 3, 4, and 6 hold. Let $\tilde{\Theta}$ be a stationary point of MSGD-P and let $\tilde{\theta}_i$ be the parameter for a probing learner $i \in U$. Then for any $\kappa \in (0, 1)$, with probability at least $1 - \kappa$ over the probing sample,*

$$\begin{aligned} \mathcal{R}(\tilde{\theta}_i) \leq & \left(1 + \frac{1}{p}\right)\epsilon + \frac{\lambda}{2}\|\theta^*\|_F^2 + R(\|\theta^*\|_F + B_\theta)\sqrt{2B_{\text{CE}}} + 4R\sqrt{\frac{(1+p)K \log K}{\lambda p n}} \\ & + (\Gamma_\star + \Gamma_B + R(\|\theta^*\|_F + B_\theta))\sqrt{\frac{\log(3/\kappa)}{2n}}. \end{aligned}$$

Proof. We work on the event $\mathcal{E}_{\text{pinsker}} \cap \mathcal{E}_\star \cap \mathcal{E}_u$ where:

- $\mathcal{E}_{\text{pinsker}}$ is the event of Lemma 25 (probability $\geq 1 - \kappa/3$),
- $\mathcal{E}_\star$ is the event of Lemma 26 (probability $\geq 1 - \kappa/3$),
- $\mathcal{E}_u$ is the event of Lemma 27 (probability $\geq 1 - \kappa/3$).

By a union bound,

$$\mathbb{P}(\mathcal{E}_{\text{pinsker}} \cap \mathcal{E}_\star \cap \mathcal{E}_u) \geq 1 - \kappa.$$

Bound on $\hat{L}_{\text{probe}}(\tilde{\theta}_i)$. By Lemma 23, dividing by p ,

$$\hat{L}_{\text{probe}}(\tilde{\theta}_i) \leq \frac{\epsilon}{p} + \hat{L}_{\text{probe}}(\theta^*) + \frac{\lambda}{2}\|\theta^*\|_F^2.$$

By Corollary 29 applied at $W = \theta^*$ and Lemma 26,

$$\hat{L}_{\text{probe}}(\theta^*) \leq \hat{L}_{\text{true}}(\theta^*) + R\|\theta^*\|_F \Delta_{1,n} \leq \epsilon + \Gamma_\star \sqrt{\frac{\log(3/\kappa)}{2n}} + R\|\theta^*\|_F \Delta_{1,n}.$$

Hence

$$\hat{L}_{\text{probe}}(\tilde{\theta}_i) \leq \frac{\epsilon}{p} + \epsilon + \frac{\lambda}{2}\|\theta^*\|_F^2 + \Gamma_\star \sqrt{\frac{\log(3/\kappa)}{2n}} + R\|\theta^*\|_F \Delta_{1,n}.$$

Bound on $\hat{L}_{\text{true}}(\tilde{\theta}_i)$. By Corollary 29 at $W = \tilde{\theta}_i$,

$$\hat{L}_{\text{true}}(\tilde{\theta}_i) \leq \hat{L}_{\text{probe}}(\tilde{\theta}_i) + R\|\tilde{\theta}_i\|_F \Delta_{1,n}.$$

Substituting the bound from the previous step,

$$\hat{L}_{\text{true}}(\tilde{\theta}_i) \leq \frac{\epsilon}{p} + \epsilon + \frac{\lambda}{2}\|\theta^*\|_F^2 + \Gamma_\star \sqrt{\frac{\log(3/\kappa)}{2n}} + R(\|\theta^*\|_F + \|\tilde{\theta}_i\|_F) \Delta_{1,n}.$$

On $\mathcal{E}_{\text{pinsker}}$, by Lemma 25,

$$\Delta_{1,n} \leq \sqrt{2B_{\text{CE}}} + \sqrt{\frac{\log(3/\kappa)}{2n}}.$$

Therefore,

$$\hat{L}_{\text{true}}(\tilde{\theta}_i) \leq \frac{\epsilon}{p} + \epsilon + \frac{\lambda}{2}\|\theta^*\|_F^2 + R(\|\theta^*\|_F + B_\theta)\sqrt{2B_{\text{CE}}} + (\Gamma_\star + R(\|\theta^*\|_F + B_\theta))\sqrt{\frac{\log(3/\kappa)}{2n}}.$$

Bound on $\mathcal{R}(\tilde{\theta}_i)$. By Lemma 27, on $\mathcal{E}_u$ and since $\|\tilde{\theta}_i\| \leq B_\theta$ by Lemma 24,

$$\mathcal{R}(\tilde{\theta}_i) \leq \hat{L}_{\text{true}}(\tilde{\theta}_i) + 2\sqrt{2}RB_\theta\sqrt{\frac{K}{n}} + \Gamma_B\sqrt{\frac{\log(3/\kappa)}{2n}}.$$

Combining with the previous bound and substituting $B_\theta = \sqrt{2(1+p)\log K/(\lambda p)}$ (which yields $2\sqrt{2}RB_\theta\sqrt{K/n} = 4R\sqrt{(1+p)(K\log K)/(\lambda pn)}$) gives the stated bound. $\square$

Lemma 23 (Stationarity bound). *At a stationary point $\tilde{\Theta}$, for learner $i \in U$,*

$$p\hat{L}_{\text{probe}}(\tilde{\theta}_i) \leq \epsilon + p\hat{L}_{\text{probe}}(\theta^\star) + \frac{\lambda p}{2}\|\theta^\star\|_F^2.$$

Proof of Lemma 23. By optimality of $\tilde{\theta}_i$ for $\Phi_i(\cdot; \tilde{\Theta})$,

$$\Phi_i(\tilde{\theta}_i; \tilde{\Theta}) \leq \Phi_i(\theta^\star; \tilde{\Theta}).$$

Subtracting the two population CE terms at $W = \tilde{\theta}_i$ and $W = \theta^\star$, and using that

$$\tau\alpha_i \mathbb{E}_{\mathcal{P}_i} \ell_{\text{CE}}(\theta^\star) + (1-\tau)a_i(\tilde{\Theta}) \mathbb{E}_{\mathcal{D}_i(\tilde{\Theta})} \ell_{\text{CE}}(\theta^\star) \leq \epsilon,$$

we obtain the stated bound. $\square$

Lemma 24 (Norm bound and probability floor). *We have $\|\tilde{\theta}_i\|_F \leq B_\theta = \sqrt{2(1+p)\log K/(\lambda p)}$. Consequently, for any x ,*

$$\min_{c \in [K]} q_{\tilde{\theta}_i}(x)_c \geq \gamma_B = \frac{1}{K} \exp(-2RB_\theta), \quad \Gamma_B = \log \frac{1}{\gamma_B} = \log K + 2RB_\theta.$$

Proof of Lemma 24. Compare $\Phi_i(\tilde{\theta}_i; \tilde{\Theta})$ to $\Phi_i(0; \tilde{\Theta})$. For $W = 0$, q_W is uniform and each CE term equals $\log K$. Thus

$$\frac{\lambda p}{2} \|\tilde{\theta}_i\|_F^2 \leq (1+p) \log K \quad \Rightarrow \quad \|\tilde{\theta}_i\|_F \leq \sqrt{\frac{2(1+p)\log K}{\lambda p}}.$$

For the floor, write for any c, c' :

$$|z_c - z_{c'}| = |(W_c - W_{c'})^\top x| \leq \|W_c - W_{c'}\| \|x\| \leq 2\|W\|_F \|x\| \leq 2R\|W\|_F.$$

Hence $q_W(x)_c \geq \frac{1}{K} \exp(-2R\|W\|_F)$. $\square$

Lemma 25 (Pseudo-label discrepancy). *Let $\Delta_{1,n} = \frac{1}{n} \sum_{q=1}^n \|\tilde{y}_i^q - y_i^q\|_1$. Under Assumption 6, for any $\kappa \in (0, 1)$, with probability at least $1 - \kappa/3$,*

$$\Delta_{1,n} \leq \sqrt{2B_{\text{CE}}} + \sqrt{\frac{\log(3/\kappa)}{2n}}.$$

Proof of Lemma 25. For one-hot y , $\text{CE}(y, \tilde{y}) = \text{KL}(y\| \tilde{y})$ and Pinsker gives $\mathbb{E} \|y - \tilde{y}\|_1 \leq \sqrt{2 \mathbb{E} \text{KL}(y\| \tilde{y})} \leq \sqrt{2B_{\text{CE}}}$. Since $\|y - \tilde{y}\|_1 \in [0, 2]$, Hoeffding yields the stated bound. $\square$

Lemma 26 (Concentration at θ^*). *With probability at least $1 - \kappa/3$,*

$$|\widehat{L}_{\text{true}}(\theta^*) - \epsilon| \leq \Gamma_\star \sqrt{\frac{\log(3/\kappa)}{2n}},$$

where $\Gamma_\star = \log K + 2R\|\theta^*\|_F$.

Proof of Lemma 26. By Lemma 24 applied to $W = \theta^*$, $\min_c q_{\theta^*}(x)_c \geq \gamma_\star$, hence $\ell_{\text{CE}}((x, y), \theta^*) \in [0, \Gamma_\star]$. Hoeffding's inequality gives the result. $\square$

Lemma 27 (Uniform convergence). *With probability at least $1 - \kappa/3$,*

$$\sup_{\|W\|_F \leq B_\theta} |\widehat{L}_{\text{true}}(W) - \mathcal{R}(W)| \leq 2\sqrt{2} R B_\theta \sqrt{\frac{K}{n}} + \Gamma_B \sqrt{\frac{\log(3/\kappa)}{2n}}.$$

Proof of Lemma 27. For any sample $\{(x_i, y_i)\}_{i=1}^n$ with $\|x_i\| \leq R$, we bound the empirical Rademacher complexity of

$$\mathcal{F}_{B_\theta} := \{(x, y) \mapsto \ell_{\text{CE}}((x, y), W) : \|W\|_F \leq B_\theta\}.$$

The function $f(z, y) = -\log \text{softmax}(z)_Y$ has gradient $\nabla_z f = q - y$ with $\|\nabla_z f\|_2 = \|q - y\|_2 \leq \sqrt{2}$. Thus f is $\sqrt{2}$ -Lipschitz in z . By the vector contraction lemma,

$$\mathfrak{R}_n(\mathcal{F}_{B_\theta}) \leq \sqrt{2} \cdot \mathbb{E} \left[\sup_{\|W\|_F \leq B_\theta} \frac{1}{n} \sum_{i=1}^n \sum_{k=1}^K \sigma_{i,k} (W x_i)_k \right],$$

with $\sigma_{i,k}$ i.i.d. Rademacher. The inner supremum equals

$$\frac{B_\theta}{n} \mathbb{E} \left\| \sum_{i=1}^n \sum_{k=1}^K \sigma_{i,k} e_k x_i^\top \right\|_F \leq \frac{B_\theta}{n} \sqrt{\mathbb{E} \left\| \sum_{i,k} \sigma_{i,k} e_k x_i^\top \right\|_F^2} = \frac{B_\theta}{n} \sqrt{\sum_{k=1}^K \sum_{i=1}^n \|x_i\|^2} \leq B_\theta R \sqrt{\frac{K}{n}}.$$

Hence $\mathfrak{R}_n(\mathcal{F}_{B_\theta}) \leq \sqrt{2} B_\theta R \sqrt{K/n}$. A standard symmetrization and bounded-difference argument (the per-sample loss range over $\|W\| \leq B_\theta$ is at most Γ_B by Lemma 24) yields, with prob. $\geq 1 - \kappa/3$,

$$\sup_{\|W\|_F \leq B_\theta} |\widehat{L}_{\text{true}}(W) - \mathcal{R}(W)| \leq 2\mathfrak{R}_n(\mathcal{F}_{B_\theta}) + \Gamma_B \sqrt{\frac{\log(3/\kappa)}{2n}},$$

which gives the stated bound. $\square$

Lemma 28 (Logits bridge for cross-entropy; no probability floors). *Let $W \in \mathbb{R}^{K \times d}$, $z(x) = Wx \in \mathbb{R}^K$, and $q_W(x) = \text{softmax}(z(x))$. For any $x \in \mathbb{R}^d$, any one-hot label $y \in \{e_1, \dots, e_K\}$, and any soft pseudo-label $\tilde{y} \in \Delta^{K-1}$ (i.e., $\tilde{y}_c \geq 0$, $\sum_c \tilde{y}_c = 1$), the cross-entropy difference satisfies the exact identity*

$$\text{CE}(\tilde{y}, q_W) - \text{CE}(y, q_W) = -\langle \tilde{y} - y, z(x) \rangle.$$

Consequently,

$$|\text{CE}(\tilde{y}, q_W) - \text{CE}(y, q_W)| \leq \|\tilde{y} - y\|_1 \cdot \|z(x)\|_\infty \leq \|\tilde{y} - y\|_1 \cdot R \|W\|_F.$$

Proof. Recall that $\text{CE}(r, q) = -\sum_{c=1}^K r_c \log q_c$ for any probability vector r , and $q_c = \frac{e^{z_c}}{\sum_j e^{z_j}}$ with $z = Wx$. Hence

$$-\log q_c = \log \left(\sum_{j=1}^K e^{z_j} \right) - z_c = \text{logsumexp}(z) - z_c.$$

Therefore, for any $r \in \Delta^{K-1}$,

$$\text{CE}(r, q_W) = \sum_{c=1}^K r_c (\text{logsumexp}(z) - z_c) = \text{logsumexp}(z) \cdot \underbrace{\sum_{c=1}^K r_c}_{=1} - \sum_{c=1}^K r_c z_c = \text{logsumexp}(z) - \langle r, z \rangle.$$

Applying this twice with $r = \tilde{y}$ and $r = y$ gives

$$\text{CE}(\tilde{y}, q_W) - \text{CE}(y, q_W) = (\text{logsumexp}(z) - \langle \tilde{y}, z \rangle) - (\text{logsumexp}(z) - \langle y, z \rangle) = -\langle \tilde{y} - y, z \rangle,$$

which is the claimed identity.

For the inequalities, use Hölder and the fact that $\|z\|_\infty = \max_c |z_c|$:

$$|\text{CE}(\tilde{y}, q_W) - \text{CE}(y, q_W)| = |\langle \tilde{y} - y, z \rangle| \leq \|\tilde{y} - y\|_1 \|z\|_\infty.$$

Finally, since $z_c = W_c^\top x$ and $\|x\| \leq R$,

$$\|z\|_\infty = \max_c |W_c^\top x| \leq \max_c \|W_c\|_2 \cdot \|x\| \leq \left(\max_c \|W_c\|_2 \right) R \leq R \|W\|_F,$$

because $\|W\|_F^2 = \sum_c \|W_c\|_2^2 \geq \max_c \|W_c\|_2^2$. Combining the bounds yields the result. $\square$

Lemma 29 (Empirical bridge on the probing batch). *On the probing dataset $\{(\tilde{x}_i^q, \tilde{y}_i^q, y_i^q)\}_{q=1}^n$, define*

$$\hat{L}_{\text{probe}}(W) = \frac{1}{n} \sum_{q=1}^n \text{CE}(\tilde{y}_i^q, q_W(\tilde{x}_i^q)), \quad \hat{L}_{\text{true}}(W) = \frac{1}{n} \sum_{q=1}^n \text{CE}(y_i^q, q_W(\tilde{x}_i^q)), \quad \Delta_{1,n} = \frac{1}{n} \sum_{q=1}^n \|\tilde{y}_i^q - y_i^q\|_1.$$

Then for any W ,

$$|\hat{L}_{\text{probe}}(W) - \hat{L}_{\text{true}}(W)| \leq R \|W\|_F \cdot \Delta_{1,n}.$$

Proof. Apply Lemma 28 termwise and average. $\square$

G Experimental Details and Additional Results

G.1 Dataset and Preprocessing Details

MovieLens-10M [21]. This dataset contains 10 million movie ratings from 70k users across 10k movies, providing a natural testbed for multi-learner competition in recommendation. Following Bose et al. [9] and Su and Dean [47], we extract $d = 16$ dimensional user embeddings via matrix factorization and retain ratings for the top 200 most-rated movies, yielding a population of 69,474 users. Each user's data consists of $z = (x, r)$ where $x \in \mathbb{R}^d$ is the embedding and r contains their ratings. Let Ω_x denote the set of movies rated by user x with $|\Omega_x|$ movies. Each learner fits a linear model $\theta \in \mathbb{R}^{d \times 200}$ using squared loss:

$$\ell(z; \theta) = \frac{1}{|\Omega_x|} \sum_{i \in \Omega_x} (\theta_i^\top x - r_i)^2.$$

ACS Employment [15]. We use the ACSEmployment task from `folktables`, where the goal is to predict employment status from demographic features. The population consists of 38,221 individuals from the 2018 Alabama census (ages 16–90), with $d = 16$ features describing age, education, marital status, etc. Each user’s data is $z = (x, y)$ where $x \in \mathbb{R}^d$ (standardized to zero mean, unit variance) and $y \in \{0, 1\}$. Each learner uses logistic regression:

$$\ell(z; \theta) = -y \log(\sigma(\theta^\top x)) - (1 - y) \log(1 - \sigma(\theta^\top x)),$$

where σ is the sigmoid function. The model predicts $\hat{y} = \mathbf{1}[\theta^\top x > 0]$.

Amazon Reviews 2023. We use the McAuley-Lab/Amazon-Reviews-2023 corpus (via Hugging-Face), constructing a binary sentiment task from review text and star ratings. We sample up to 30,000 reviews from nine product categories and define labels by $y = \mathbf{1}[\text{rating} \geq 4]$ (so 1–3 stars are negative, 4–5 stars are positive). Features are $d = 384$ dimensional sentence embeddings of the review text produced by `all-MiniLM-L6-v2`, with a stratified 95/5 train-test split. Each learner uses the same logistic regression setup as Census.

User Preferences. For Census and MovieLens, we simulate a market with $m = 5$ learners and model inherent user preferences $\pi(z)$ via K-means clustering ($K = 5$) on user features, assigning each user’s preferred platform based on their cluster membership. For Amazon, we use category-based partitions ($m = 9$ learners, one per product category): each review is assigned to its category group, and this group index induces the preference partition $\{S_i\}_{i=1}^m$. In all cases, the partition captures the intuition that users from different demographic or behavioral segments may have systematic affinities for different platforms.

Train/test splits. Each experiment constructs one held-out test set at the start and uses the remaining users as the shared training pool for all learners. Census uses a fixed 99/1 split with random seed 0; Amazon uses a stratified 95/5 split over sentiment labels; MovieLens uses a 90/10 user split. We do not use a validation split because the experimental hyperparameters are fixed in advance rather than tuned per run. Dataset-specific hyperparameters are summarized in Table 3.

Census neural experiment. The neural Census experiment in Figure 6 changes only the learner family relative to the linear Census experiments. Each learner is a two-layer ReLU MLP with architecture $d \rightarrow 64 \rightarrow 1$, trained with SGD on binary cross-entropy. The held-out test split, K-means preference clusters, and cluster-induced learner rankings are reused unchanged. In the bad-outcome run, learners are randomly initialized, $\tau = 0.3$, and no learner probes. In the good-outcome run, learners are initialized by partition pretraining on the users who rank them first, $\tau = 0.7$, and Learner 2 probes. Probing is offline: pseudo-labels are collected once from the initialized peer selected by the preference-aware rule $T_i(x) = \{\pi(x)\}$, with $\kappa = 0$, and the probing learner then trains on the weighted combination of organic and pseudo-labeled losses. We use $T = 1000$ MSGD iterations, three random seeds, and the probing weights shown in Figure 6.

G.2 Additional Experiments

In this section, we provide additional experiments, which elaborate on Section 6.

Parameter	Description	Census	MovieLens	Amazon
m	Number of learners	5	5	9
T	Total rounds	4000	4000	20000
λ	L2 regularization	10^{-9}	10^{-3}	10^{-9}
n	Offline probe dataset size	100	1000	500

Table 3: Hyperparameters by dataset. Shared values are merged across columns.

Expt 1 for other scenarios We replicate Expt 1 from Section 6 in the two globally-good settings from Definition 2: market-leader and majority-good. As in the preference-aware case, standard MSGD without probing ($p = 0$, $\tau = 0.3$) converges to equilibria with large full-population gaps to the dashed black baseline (Figures 8 and 10). In the market-leader setting (Figure 8), the most overspecialized learner is about 0.32 below baseline on Census (≈ 0.47 vs ≈ 0.79), and more than 2.6 MSE above baseline on MovieLens (> 5.6 vs ≈ 2.95). In the majority-good setting (Figure 10), two learners remain overspecialized with sizable baseline gaps (Census around 0.22–0.25 below baseline; MovieLens around 2.0–2.9 above baseline). These results show that poor global equilibria are not unique to the preference-aware scenario.

Expt 2 for other scenarios We next replicate Expt 2 for the other scenarios. In the market-leader scenario, Learner 4 probes the known leader (Learner 0), and its final Census accuracy improves from about 0.55 to about 0.75, while its MovieLens loss drops from about 5.1 to about 3.1 as p increases (Figure 9). Relative to the dashed baseline, this closes most of the initial gap (roughly from 0.24 to 0.04 on Census, and from about 2.2 to about 0.1 on MovieLens). In the majority-good scenario, where Learner 4 probes via median aggregation over all peers, we observe a similar pattern: Census accuracy rises from about 0.52 to about 0.77, and MovieLens loss decreases from about 4.7 to about 3.1 (Figure 11). This again closes a large fraction of the baseline gap (roughly from 0.27 to 0.02 on Census, and from about 1.7 to about 0.1 on MovieLens). Well-performing learners change only modestly, indicating that probing primarily benefits the underperforming learner and mitigates overspecialization.

Expt 4: Impact of noise in selection of probed labels Finally, we test robustness to noisy probing-source selection. For each probe query x , the probing learner queries $\pi(x)$ with probability $1 - \kappa$, and with probability κ it queries a random other learner. Across Census, Amazon, and MovieLens (Figure 4), increasing κ causes only mild changes in the probing learner’s final performance relative to the low-noise case, while preserving strong gains from probing. Thus, our method is robust to imperfect estimates of the ranking function.

Expt 5: What happens when multiple learners probe? We also evaluate a preference-aware setting where multiple learners probe simultaneously. In Figure 12, the triangle-marked learners (Learners 2 and 3) both probe while the dashed black line indicates the full-data baseline. As p increases, Learner 2 improves from about 0.60 to about 0.78, and Learner 3 improves from about 0.66 to about 0.79. Equivalently, their baseline gaps shrink from roughly 0.19 and 0.13 at $p = 0$ to about 0.01 and near zero at $p = 0.8$. The non-probing learners move only slightly, indicating that simultaneous probing remains stable and still helps underperforming learners recover most of the overspecialization gap.

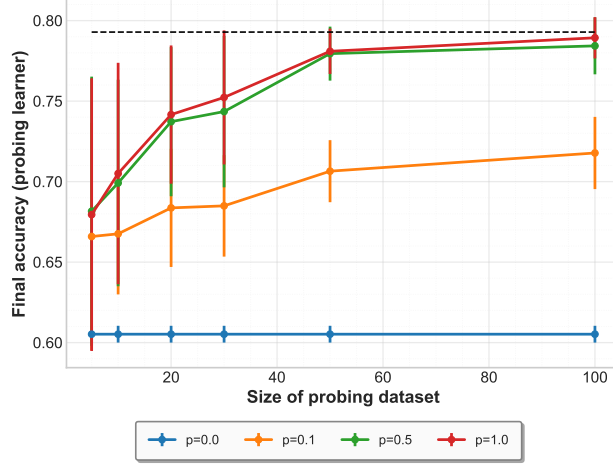


Figure 7: **Performance of probing learner on Census as a function of n .** Error bars show one standard deviation over 10 random seeds.

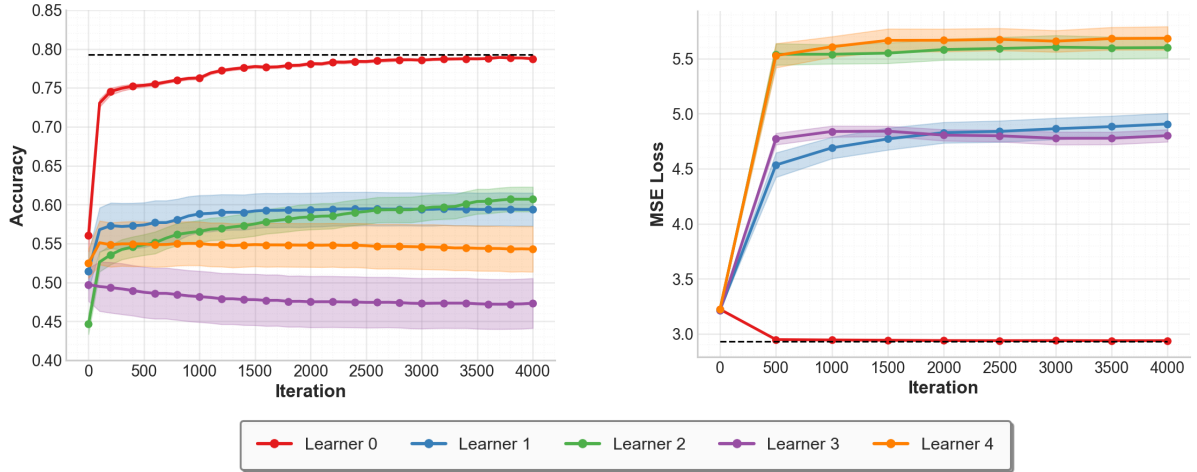


Figure 8: **MSGD full-population performance with random initialization (Market-leader scenario).** Left: Census test accuracy. Right: MovieLens test loss. Here $\tau = 0.3$.

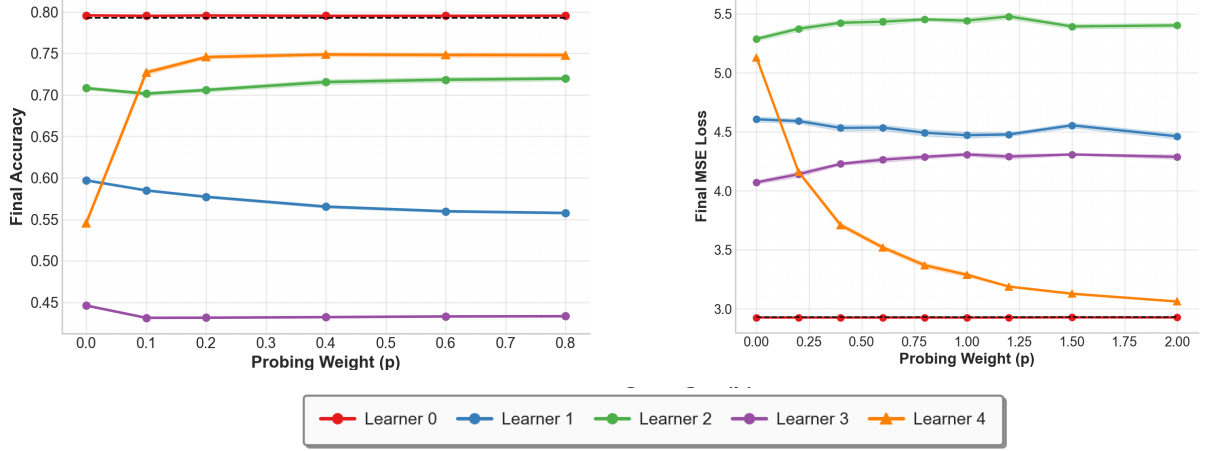


Figure 9: **Effect of probing on full-population performance (Market-leader scenario)** Left: Final accuracy vs probing weight p on Census. Right: MovieLens final loss vs p . The triangle markers indicate Learner 4 probes the market leader, learner 0. Here $\tau = 0.7$.

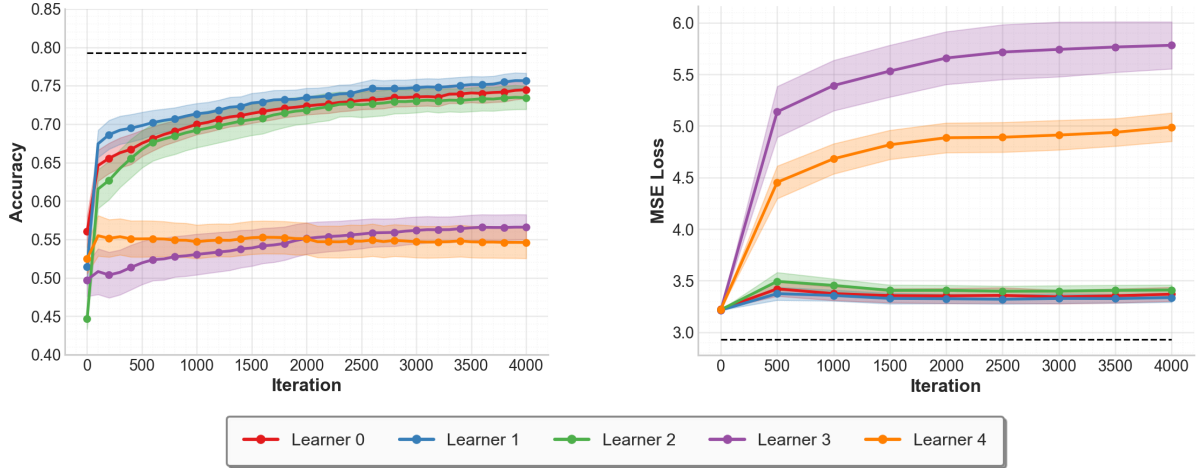


Figure 10: **MSGD full-population performance with random initialization (Majority good scenario)**. Left: Census test accuracy. Right: MovieLens test loss. Here $\tau = 0.3$.

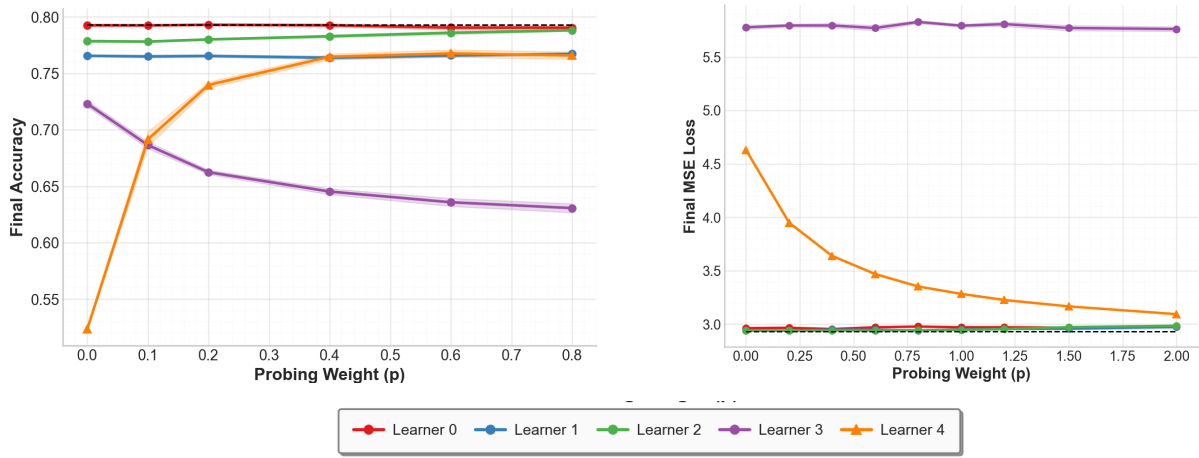


Figure 11: **Effect of probing on full-population performance (Majority Good Scenario)** Left: Final accuracy vs probing weight p on Census. Right: MovieLens final loss vs p . Triangle markers indicate Learner 4 is probing via median aggregation over all peers. Here $\tau = 0.7$.

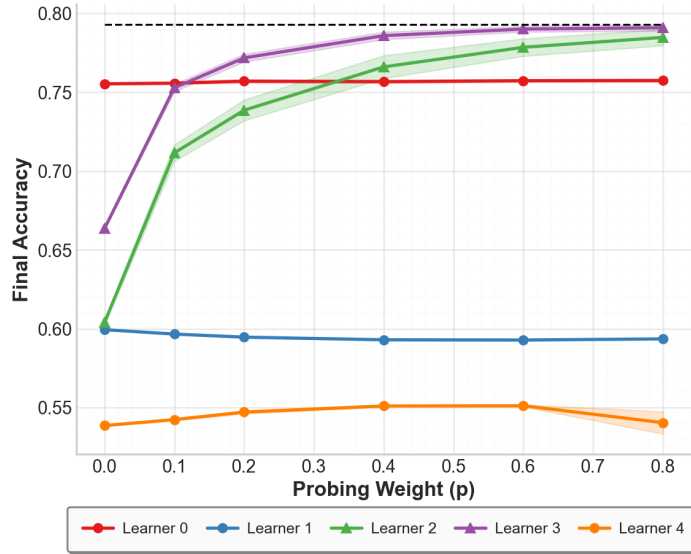


Figure 12: **Effect of probing on full-population performance when multiple learners probe (Preference-aware scenario)**. Census final accuracy vs probing weight p . Triangle markers indicate the probing learners (Learners 2 and 3). The dashed black line denotes the full-data baseline.