

Open Problem: The Sample Complexity of Multi-Distribution Learning for VC Classes

Pranjal Awasthi

Google Research, Mountain View, CA, USA

PRANJALAWASTHI@GOOGLE.COM

Nika Haghtalab

University of California, Berkeley, CA, USA

NIKA@BERKELEY.EDU

Eric Zhao

University of California, Berkeley, CA, USA

ERIC.ZH@BERKELEY.EDU

Editors: Gergely Neu and Lorenzo Rosasco

Abstract

Multi-distribution learning is a natural generalization of PAC learning to settings with multiple data distributions. There remains a significant gap between the known upper and lower bounds for PAC-learnable classes. In particular, though we understand the sample complexity of learning a VC dimension d class on k distributions to be $O(\varepsilon^{-2} \ln(k)(d + k) + \min \{\varepsilon^{-1} dk, \varepsilon^{-4} \ln(k)d\})$, the best lower bound is $\Omega(\varepsilon^{-2}(d + k \ln(k)))$. We discuss recent progress on this problem and some hurdles that are fundamental to the use of game dynamics in statistical learning.

Keywords: PAC learning, multi-distribution learning, distributional robustness, learning in games.

1. Introduction

The pervasive need for robustness, fairness, and multi-agent welfare in learning processes has led to the development of learning paradigms whose performance hold under multiple distributions and scenarios. *Multi-distribution learning*, or MDL, is a setting introduced by [HJZ22] to address these needs and unify several existing frameworks and applications, such as notions of *min-max* fairness [MSS19, AAK⁺22], *group distributionally robust* optimization [SKHL20], and collaborative learning [BHPQ17]. MDL is a generalization of the agnostic learning paradigms [Val84, BEHW89] to multiple data distributions. In this setting, given a set of distributions $\mathcal{D} = \{D_1, \dots, D_k\}$ supported on $\mathcal{X} \times \mathcal{Y}$, loss function ℓ , and a hypothesis class $\mathcal{H}$, the goal of MDL is to find a (possibly randomized) hypothesis h where

$$\max_{D \in \mathcal{D}} \mathcal{L}_D(h) \leq \varepsilon + \min_{h^* \in \mathcal{H}} \max_{D \in \mathcal{D}} \mathcal{L}_D(h^*), \text{ where } \mathcal{L}_D(h) := \mathbb{E}_{(x,y) \sim D} [\ell(h, (x, y))]. \quad (1)$$

Such an h is called an ε -optimal solution to the MDL problem $(\mathcal{D}, \mathcal{H})$ and we denote $\text{OPT} := \min_{h^* \in \mathcal{H}} \max_{D \in \mathcal{D}} \mathcal{L}_D(h^*)$. Our open problem concerns the sample complexity of MDL.

Problem Statement. Consider an example oracle EX_i for each distribution $D_i \in \mathcal{D}$, which once queried returns an independent sample $(x, y) \sim D_i$. The optimal sample complexity of MDL is the smallest total number of queries issued to examples oracles, in a possibly adaptive fashion, that is sufficient for learning an ε -optimal solution. Formally, a multi-distribution learning algorithm at each iteration $t = 1, 2, \dots$, chooses an index $i^{(t)} \in [k]$, queries $\text{EX}_{i^{(t)}}$ to sample an instance $(x^{(t)}, y^{(t)})$ and, upon termination, returns a (possibly randomized) solution h . We use the shorthands $z^{(t)} = (x^{(t)}, y^{(t)}, i^{(t)})$, $\mathcal{Z} = \mathcal{X} \times \mathcal{Y} \times [k]$, and $\mathcal{Z}^*$ to denote a sequence $z^{(1)}, z^{(2)}, \dots$ of any size.

Definition 1 (Multi-Distribution Learnability) We say a hypothesis class $\mathcal{H}$ is multi-distribution learnable with sample complexity $m_{\mathcal{H}} : (0, 1)^2 \times \mathbb{N} \rightarrow \mathbb{N}$ if there exists functions $\mathcal{A}_s : \mathcal{Z}^* \rightarrow [k]$ and $\mathcal{A}_h : \mathcal{Z}^* \rightarrow \Delta(\mathcal{Y})^{\mathcal{X}}$ where the following holds: for every $(\varepsilon, \delta) \in (0, 1)$, $k \in \mathbb{N}$, and set of k distributions $\mathcal{D}$ over $\mathcal{X} \times \mathcal{Y}$, by letting $i^{(t)} = \mathcal{A}_s(z^{(1)}, \dots, z^{(t-1)})$ for $t \in [m_{\mathcal{H}}(\varepsilon, \delta, k)]$, with probability at least $1 - \delta$, the solution $h = \mathcal{A}_h(z^{(1)}, \dots, z^{(m)})$ is ε -optimal, i.e., satisfying (1).

Problem 1 What is the optimal sample complexity of MDL? Are hypothesis classes $\mathcal{H}$ with VC dimension d multi-distribution learnable with a sample complexity of $O(\varepsilon^{-2}(\ln(k)d + k \ln(k/\delta)))$?

Recalling that the sample complexity of agnostic learning is $m_{\mathcal{H}}(\varepsilon, \delta, 1) \in \Theta(\varepsilon^{-2}(d + \ln(1/\delta)))$ [SB14], one hopes to avoid paying the $\Omega(k \cdot m_{\mathcal{H}}(\varepsilon, \delta/k, 1))$ samples necessary to independently learn each of the k data distributions. This is why our conjectured sample complexity avoids a dependence on dk and has an optimal ε^{-2} dependence. Existing results, however, have fallen short of meeting both of these requirements and traded off lack of dependence on dk with the optimal dependence on ε , as shown in rows 1 and 2 of Table 1. On the other hand, the optimal sample complexity of MDL has been rightly characterized for finite hypothesis classes in row 3 (and more generally those of finite Littlestone dimension or Bregman diameter [HJZ22]) and obtains optimal $\varepsilon^{-2} \ln(|\mathcal{H}|)$ dependence. The best lower bound, row 4, leaves a logarithmic gap with the conjectured upper bound. Near-optimal bounds are known for *realizable* settings where $\text{OPT} = 0$ (row 5) and *personalized* settings where one can produce a different hypothesis for each distribution (row 6).

Table 1: Best known bounds on the sample complexity of MDL for hypothesis classes with VC dimension d . $\tilde{O}$ hides double-log factors and an additive factor of $\varepsilon^{-2}k \ln(k/\delta)$.

Bound	Assumption	Citation
1. $\tilde{O}(\varepsilon^{-2} \ln(k)d + \varepsilon^{-1} dk \log(d/\varepsilon))$	N/A	[HJZ22]
2. $\tilde{O}(\varepsilon^{-4} \ln(k)(d + \ln(1/\delta\varepsilon)))$	N/A	(Theorem 7)
3. $\tilde{O}(\varepsilon^{-2} \ln(\mathcal{H}))$	N/A	[HJZ22]
4. $\Omega(\varepsilon^{-2}(d + k \ln(\min\{d, k\}/\delta)))$	N/A	[HJZ22]
5. $O(\ln(k)\varepsilon^{-1}(d \ln(1/\varepsilon) + k \ln(k/\delta)))$	$\text{OPT} = 0$	[CZZ18, NZ18]
6. $\tilde{O}(\ln(k)\varepsilon^{-2}(d \ln(d/\varepsilon) + k \ln(k/\delta)))$	Personalized	(Theorem 9)

Broad Applications. One of the motivating application of MDL is *collaborative learning*, where multiple stakeholders (representing D_i) collaborate in training a model that provides high performance for each stakeholder [BHPQ17, NZ18, CZZ18, BHPS21]. The sample complexity of MDL thus quantifies the value of collaboration in learning: whereas our conjectured upper bound would imply that collaboration reduces the amount of data needed by a $\ln(k)/k$ factor, existing bounds only imply a $\min\{\ln(k)/k\varepsilon^2, \varepsilon\}$ factor reduction.

Another application of MDL is to Group *distributionally robust optimization* (DRO) which concerns learning a model with performance guarantees for many deployment environments [SKHL20, SRKL20]. MDL sample complexity bounds quantify the cost of obtaining this robustness, a question of growing interest and which has been studied in terms of finite-sum convergence [CH22,

ACJ⁺21] and sample complexity [HJZ22]. Our conjectured upper bound would extend these favorable results to VC classes by only increasing the sample complexity logarithmically.

MDL also captures notions of min-max fairness in learning, which concerns prioritizing the well-being of the worst-off subgroup and has applications in federated learning [MSS19] and equity [AAK⁺22]. Min-max fair learning has mainly been studied in settings with presampled datasets, where an inevitable sample complexity lower bound of $\Omega(dk/\varepsilon^2)$ arises as one cannot adaptively choose distributions to sample from. The sample complexity of MDL thus captures how min-max fairness can be attained at less cost by adapting one’s data collection strategy on the fly.

2. Overview of Current Approaches

Multi-distribution learning can be formulated as the zero-sum game between a “learner” who chooses hypotheses $h \in \mathcal{H}$ and an “adversary” who chooses indices $i \in [k]$, with the payoff function $\mathcal{L}_{D_i}(h)$. Importantly, for any mixed-strategy ε -min-max equilibrium $(p, q) \in \Delta(\mathcal{H}) \times \Delta_k$, the randomized map p is a 2ε -optimal solution. All existing multi-distribution learning algorithms can be expressed as finding a ε -equilibrium using no-regret dynamics (see [HJZ22] for an overview).

Game dynamics. Formally, a game dynamic is a T -iteration process where, at each $t \in [T]$, a learner chooses hypothesis $h^{(t)} \in \mathcal{H}$ with a no-regret algorithm and an adversary chooses a distribution $i^{(t)} \in [k]$ with a (semi-)bandit algorithm. The learner estimates its current cost function $h \mapsto \mathcal{L}_{D_{i^{(t)}}}(h)$ by sampling N_{learn} datapoints from $\text{EX}_{i^{(t)}}$, while the adversary estimates its cost function $i \mapsto -\mathcal{L}_{D_i}(h^{(t)})$ by, for N_{adv} choices of $i \in [k]$, sampling a datapoint from each EX_i . The random mapping p where $p(x) = \text{Uniform}(h^{(1)}(x), \dots, h^{(t)}(x))$ is a 2ε -optimal solution.

Different instantiations. Every result in Table 1 can be obtained by instantiating this game dynamics template. Row 3 can be obtained by setting $N_{\text{learn}} = N_{\text{adv}} = 1$, $T \propto \varepsilon^{-2}(\ln(|\mathcal{H}|) + k \ln(k/\delta))$, having the learner choose $h^{(t)}$ with Hedge and the adversary choose $i^{(t)}$ with Exp3 [HJZ22]. Row 1 can be obtained with the same algorithm but first creating an offline ε -covering of the class $\mathcal{H}$ on each data distribution $D_i \in \mathcal{D}$, using $O(d/\varepsilon)$ samples per distribution. Row 2 can be obtained by setting $N_{\text{adv}} = k$, $N_{\text{learn}} \propto \varepsilon^{-2}(d + \ln(1/\delta\varepsilon))$, $T \propto \varepsilon^{-2} \ln(k/\delta)$, having the learner choose $h^{(t)}$ to be the (approximate) risk minimizer of the current cost function and the adversary choose $i^{(t)}$ with Hedge (Theorem 7); in contrast to the prior upper bound, this bound uses an algorithm that iterates fewer times but samples more at each iteration.

Personalization. We can pinpoint the challenge of negotiating trade-offs between different data distributions as the primary difficulty of handling infinite classes. Consider the personalized setting where, during inference time, $\mathcal{A}_h(z^{(1)}, \dots, z^{(m)})$ can return a different hypothesis h_i for each distribution D_i . This assumes away the difficulty of combining hypotheses that are each near-optimal for different distributions. The conjectured sample complexity bound of $\tilde{O}(\ln(k)\varepsilon^{-2}(d \ln(d/\varepsilon) + k \ln(k/\delta)))$ can be obtained in the personalized setting (Row 6 of Table 1) by running the Row 1 algorithm $\ln(k)$ times, at each round limiting the adversary to playing within a small region of the simplex Δ_k that we can efficiently cover $\mathcal{H}$ on (Theorem 9).

2.1. Existing Challenges

Adaptive coverings. A potential approach to closing the gap with the conjectured sample complexity bound is to find a method of adaptively covering the hypothesis class $\mathcal{H}$. Whereas Row 1

was obtained by taking a naive offline ε -covering of $\mathcal{H}$ on all k distributions, Row 2 was obtained by an algorithm that (implicitly) ε -covers the class $\mathcal{H}$ on $O(\ln(k)\varepsilon^{-2})$ adaptive choices of $D_i \in \mathcal{D}$. It is unclear whether a covering of lower resolution can be used, or if it is possible to only cover $\mathcal{H}$ on $O(\ln(k))$ choices of distributions $D_i \in \mathcal{D}$. We also note that it is not the size of the ε -covering of k distributions, i.e., $k\varepsilon^{-O(d)}$, that is the bottleneck, but rather the number of samples needed to create such a cover. In contrast, the personalized algorithm decided in an online fashion what distributions need to be covered and it only covers $\mathcal{H}$ on $O(\ln(k))$ choice of (mixture) distributions from $\mathcal{D}$.

Agnostic-to-realizable. Another potential tool is an agnostic-to-realizable reduction [HKLM22], since nearly-optimal sample complexity bounds are known for realizable settings where $\text{OPT} = 0$ [BHPQ17, CZZ18, NZ18]. This technique has had success in related problems, such as the closely related adversarial PAC learning problem [MHS19]. Unfortunately, because multi-distribution learning involves online decision-making—determining which example oracles to call—the usual reduction of testing all possible labelings of observed datapoints is intractable.

Bounding regret. Game dynamics algorithms rely on the learner achieving a low regret on the sequence of distributions chosen by the adversary. However, with VC classes, even when all distributions share a Bayes classifier, an oblivious adversary can force the learner to suffer regret linear in k . It is therefore necessary to reason about the adversary’s behavior to bound the regret of the learner. This is atypical; game dynamics proofs usually bound each player’s regret independently.

Proposition 2 *Consider an algorithm $\mathcal{A}$ that, given distributions $D_1, \dots, D_T$, draws only N datapoints in total and returns a sequence of hypotheses $h_1, \dots, h_k$ where each h_t is trained only on datapoints sampled from $D_1, \dots, D_t$. There exists a sequence $D_1, \dots, D_T$ with only k distinct members, where $\mathbb{E}[T^{-1} \sum_{t \in [T]} \mathcal{L}_{D_t}(h_t)] - \min_{h^* \in \mathcal{H}} T^{-1} \sum_{t \in [T]} \mathcal{L}_{D_t}(h^*) \in \Omega(\sqrt{dk/N})$.*

3. Intermediate Open Problems

Lower Bounds. We believe a $\ln(k)d$ factor is missing from the best known sample complexity lower bound of $\Theta(\varepsilon^{-2}(d + k \ln(\min\{k, d\}/\delta)))$. The absence of a $\ln(k)d$ term would be significant as it would imply that, when VC dimension dominates sample complexity, handling more data distributions comes effectively for free. Interestingly, this $\ln(k)$ factor does not appear in the upper bound when the complexity of $\mathcal{H}$ is characterized by Littlestone dimension, perhaps due to the stronger compression guarantees for online-learnable classes. A $\ln(k)d$ term would also shed light on compression schemes for VC classes [LW86]; a lower bound of $\Theta(\ln(k)d + k)$ would lend evidence against the existence of $O(\text{VC}(\mathcal{H}))$ -size compression schemes.

Problem 2 *Is the sample complexity of multi-distribution learning in $\Omega(\log(k)d)$?*

Proper learning. All existing multi-distribution learning algorithms with fast sample complexity rates produce either a randomized hypothesis $h \in \Delta(\mathcal{H})$ or an improper hypothesis resulting from taking a majority vote. An open question is whether improperness is necessary for fast rates.

Problem 3 *What is the sample complexity of proper multi-distribution learning?*

Oracle-efficient learning. For oracle-efficient algorithms, that is an algorithm only accessing $\mathcal{H}$ through an ERM oracle [DHL⁺20], only the sample complexity bound from Row 2 in Table 1 is known. An open question is whether there exists a statistical-computational trade-off for MDL.

Problem 4 *What is the sample complexity of oracle-efficient multi-distribution learning?*

References

- [AAK⁺22] Jacob D. Abernethy, Pranjal Awasthi, Matthäus Kleindessner, Jamie Morgenstern, Chris Russell, and Jie Zhang. Active sampling for min-max fairness. In Kamalika Chaudhuri, Stefanie Jegelka, Le Song, Csaba Szepesvári, Gang Niu, and Sivan Sabato, editors, *Proceedings of the International Conference on Machine Learning (ICML)*, volume 162 of *Proceedings of Machine Learning Research*, pages 53–65. PMLR, 2022.
- [ACJ⁺21] Hilal Asi, Yair Carmon, Arun Jambulapati, Yujia Jin, and Aaron Sidford. Stochastic bias-reduced gradient methods. In Marc’Aurelio Ranzato, Alina Beygelzimer, Yann N. Dauphin, Percy Liang, and Jennifer Wortman Vaughan, editors, *Advances in Neural Information Processing Systems*, pages 10810–10822, 2021.
- [BEHW89] Anselm Blumer, Andrzej Ehrenfeucht, David Haussler, and Manfred K Warmuth. Learnability and the vapnik-chervonenkis dimension. *Journal of the ACM (JACM)*, 36(4):929–965, 1989.
- [BHPQ17] Avrim Blum, Nika Haghtalab, Ariel D. Procaccia, and Mingda Qiao. Collaborative PAC learning. In I. Guyon, U. Von Luxburg, S. Bengio, H. Wallach, R. Fergus, S. Vishwanathan, and R. Garnett, editors, *Advances in Neural Information Processing Systems 30*, pages 2392–2401. Curran Associates, Inc., 2017.
- [BHPS21] Avrim Blum, Nika Haghtalab, Richard Lanus Phillips, and Han Shao. One for one, or all for all: equilibria and optimality of collaboration in federated learning. In Marina Meila and Tong Zhang, editors, *Proceedings of the International Conference on Machine Learning (ICML)*, volume 139 of *Proceedings of Machine Learning Research*, pages 1005–1014. PMLR, 2021.
- [CH22] Yair Carmon and Danielle Hausler. Distributionally robust optimization via ball oracle acceleration. In *Advances in Neural Information Processing Systems*, 2022.
- [CZZ18] Jiecao Chen, Qin Zhang, and Yuan Zhou. Tight bounds for collaborative PAC learning via multiplicative weights. In S. Bengio, H. Wallach, H. Larochelle, K. Grauman, N. Cesa-Bianchi, and R. Garnett, editors, *Advances in Neural Information Processing Systems 31*, pages 3602–3611. Curran Associates, Inc., 2018.
- [DHL⁺20] Miroslav Dudik, Nika Haghtalab, Haipeng Luo, Robert E. Schapire, Vasilis Syrgkanis, and Jennifer Wortman Vaughan. Oracle-efficient online learning and auction design. *J. ACM*, 67(5):26:1–26:57, 2020.
- [FS97] Yoav Freund and Robert E. Schapire. A decision-theoretic generalization of on-line learning and an application to boosting. *Journal of Computer and System Sciences*, 55(1):119–139, 1997.
- [HJZ22] Nika Haghtalab, Michael I. Jordan, and Eric Zhao. On-Demand Sampling: Learning Optimally from Multiple Distributions. *CoRR*, abs/2210.12529, 2022.
- [HKLM22] Max Hopkins, Daniel M. Kane, Shachar Lovett, and Gaurav Mahajan. Realizable learning is all you need. In Po-Ling Loh and Maxim Raginsky, editors, *Proceedings of*

- the Conference on Learning Theory (COLT)*, volume 178 of *Proceedings of Machine Learning Research*, pages 3015–3069. PMLR, 2022.
- [HW86] D Haussler and E Welzl. Epsilon-nets and simplex range queries. In *Proceedings of the Second Annual Symposium on Computational Geometry*, SCG ’86, page 61–71. Association for Computing Machinery, 1986.
- [LW86] Nick Littlestone and Manfred Warmuth. Relating data compression and learnability. 1986.
- [MHS19] Omar Montasser, Steve Hanneke, and Nathan Srebro. VC classes are adversarially robustly learnable, but only improperly. In Alina Beygelzimer and Daniel Hsu, editors, *Proceedings of the Conference on Learning Theory (COLT)*, volume 99 of *Proceedings of Machine Learning Research*, pages 2512–2530. PMLR, 2019.
- [MSS19] Mehryar Mohri, Gary Sivek, and Ananda Theertha Suresh. Agnostic federated learning. In Kamalika Chaudhuri and Ruslan Salakhutdinov, editors, *Proceedings of the International Conference on Machine Learning (ICML)*, volume 97 of *Proceedings of Machine Learning Research*, pages 4615–4625. PMLR, 2019.
- [NJLS09] Arkadi Nemirovski, Anatoli Juditsky, Guanghui Lan, and Alexander Shapiro. Robust stochastic approximation approach to stochastic programming. *SIAM Journal on Optimization*, 19(4):1574–1609, 2009.
- [NZ18] Huy L. Nguyen and Lydia Zakynthinou. Improved algorithms for collaborative PAC learning. In S. Bengio, H. Wallach, H. Larochelle, K. Grauman, N. Cesa-Bianchi, and R. Garnett, editors, *Advances in Neural Information Processing Systems 31*, pages 7642–7650. Curran Associates, Inc., 2018.
- [SB14] Shai Shalev-Shwartz and Shai Ben-David. *Understanding Machine Learning - From Theory to Algorithms*. Cambridge University Press, 2014.
- [SKHL20] Shiori Sagawa, Pang Wei Koh, Tatsunori B. Hashimoto, and Percy Liang. Distributionally robust neural networks. In *Proceedings of the International Conference on Learning Representations (ICLR)*. OpenReview, 2020.
- [SRKL20] Shiori Sagawa, Aditi Raghunathan, Pang Wei Koh, and Percy Liang. An investigation of why overparameterization exacerbates spurious correlations. In Hal Daumé III and Aarti Singh, editors, *Proceedings of the International Conference on Machine Learning (ICML)*, volume 119 of *Proceedings of Machine Learning Research*, pages 8346–8356. PMLR, 2020.
- [Val84] Leslie G. Valiant. A theory of the learnable. In *Proceedings of the Annual ACM SIGACT Symposium on Theory of Computing (STOC)*, pages 436–445. ACM, 1984.

Appendix A. Omitted Proofs

We first recall standard results in online learning. We use the shorthands $x^{(1:T)} := x^{(1)}, \dots, x^{(T)}$, $\{f(x^{(t)})\}^{(1:T)} := f(x^{(1)}), \dots, f(x^{(T)})$, and $f(\cdot, b) := a \mapsto f(a, b)$ throughout this section. We use $\Delta(A)$ to denote the set of probability distributions over a set A , and Δ_d to denote a probability simplex in $\mathbb{R}^{d-1}$. Given a distribution $\mathcal{P} \in \Delta_d$, we use $(\Delta_d)_2$ to denote the convex subset of Δ_d that is the distributions that are 2-smooth: $(\Delta_d)_2 := \{\mathcal{P} \in \Delta_d \mid \max_i \mathcal{P}_i \leq 2/d\}$.

Online learning. For a sequence of actions $a^{(1)}, \dots, a^{(T)} \in A$ and costs $c^{(1)}, \dots, c^{(T)} : A \rightarrow [0, 1]$, regret is defined as $\text{Reg}(a^{(1:T)}, c^{(1:T)}) := \sum_{t=1}^T c^{(t)}(a^{(t)}) - \min_{a^* \in A} \sum_{t=1}^T c^{(t)}(a^*)$. An online learning algorithm Alg maps from costs $c^{(1:t-1)}$ to a new action $a^{(t)} \in A$, where $a^{(t)} = \text{Alg}_A(c^{(1:t-1)})$. We recall the following online learning regret bound for probability simplices.

Lemma 3 *Let $\mathcal{A}$ be a compact convex subset of Δ_d and fix a learning rate $\eta \in [0, 0.5]$. For any sequence of linear costs $c^{(1:T)}$, the Hedge online learning algorithm [FS97] chooses actions $a^{(1:T)}$, where $a^{(t)} = \text{Hedge}_{\mathcal{A}}(c^{(1:t-1)})$, with regret $\text{Reg}(a^{(1:T)}, c^{(1:T)}) \leq \ln(d)/\eta + \eta \min_{a^* \in A} \sum_{t=1}^T c^{(t)}(a^*)$.*

Stochastic costs are functions $\hat{c} : A \times \mathcal{Z} \rightarrow [0, 1]$ of both actions and datapoints. We say a stochastic cost $\hat{c}$ is linear if $\hat{c}(\cdot, z)$ is linear in its first argument under any datapoint $z \in \mathcal{Z}$. We know that estimating stochastic costs with i.i.d. samples does not significantly affect the regret of an online learning algorithm.

Lemma 4 *Let $\mathcal{A}$ be a compact convex subset of Δ_d , Alg an online learning algorithm, and $z^{(1:T)} \stackrel{\text{i.i.d.}}{\sim} D$ i.i.d. samples from some data distribution D . For any sequence of linear stochastic costs $\hat{c}^{(1:T)}$, applying Alg to the empirical cost estimates $\{\hat{c}^{(t)}(a, z^{(t)})\}^{(1:T)}$ such that $a^{(t)} = \text{Alg}_A(\{\hat{c}^{(\tau)}(a, z^{(\tau)})\}^{(1:t-1)})$ guarantees*

$$\left| \text{Reg}(a^{(1:T)}, \{\mathbb{E}_{z \sim D}[\hat{c}^{(t)}(\cdot, z)]\}^{(1:T)}) - \text{Reg}(a^{(1:T)}, \{\hat{c}^{(\tau)}(a, z^{(\tau)})\}^{(1:T)}) \right| \leq O\left(\sqrt{\ln(d/\delta)T}\right),$$

with probability at least $1 - \delta$ over the randomness of $z^{(1:T)}$ [NJLS09].

We also recall the agnostic learning upper bound.

Lemma 5 *Consider any stochastic cost $\hat{c} : A \times \mathcal{Z} \rightarrow [0, 1]$ and data distribution D , where d is the VC dimension of A . With only $O((d + \ln(1/\delta))/\varepsilon\alpha)$ samples from D , the action $a \in A$ empirically minimizing $\hat{c}$ is ε -optimal with probability $1 - \delta$: $\mathbb{E}_{z \sim D}[\hat{c}(a, z)] \leq \varepsilon + (1 + \alpha) \min_{a^* \in A} \mathbb{E}_{z \sim D}[\hat{c}(a^*, z)]$ [NZ18].*

Finally, we note that all the aforementioned results for cost sequences also apply to *payoff sequences*, where the regret of actions $a^{(1:T)}$ with respect to a sequence of payoffs $\rho^{(1:T)}$ is defined as $\text{Reg}_+(a^{(1:T)}, \rho^{(1:T)}) := \max_{a^* \in A} \sum_{t=1}^T \rho^{(t)}(a^*) - \sum_{t=1}^T \rho^{(t)}(a^{(t)})$. Here, we use the subscript $+$ in Reg_+ to distinguish when regrets are stated for payoff functions. For example, the regret bound of Hedge for payoffs can be written as follows.

Lemma 6 *Let $\mathcal{A}$ be a compact convex subset of Δ_d and fix a learning rate $\eta \in [0, 0.5]$. For any sequence of payoffs $\rho^{(1:T)}$, the Hedge online learning algorithm [FS97] chooses actions $a^{(1:T)}$, where $a^{(t)} = \text{Hedge}_{\mathcal{A}}(\rho^{(1:t-1)})$, with regret $\text{Reg}_+(a^{(1:T)}, \rho^{(1:T)}) \leq \ln(d)/\eta + \eta \max_{a^* \in A} \sum_{t=1}^T \rho^{(t)}(a^*)$.*

Algorithm 1 Multi-Distribution Learning Algorithm.

Input: Hypotheses $\mathcal{H}$, distributions $\mathcal{D}$, iterations $T \in \mathbb{Z}_+$, sub-iterations $r_1, r_2 \in \mathbb{Z}_+$, parameter $\alpha \in (0, 0.5)$;
 Initialize Hedge iterate $D^{(1)}$ to be a uniform mixture of $\mathcal{D}$;
for $t = 1, 2, \dots, T$ **do**
 Sample r_1 datapoints $z_1, \dots, z_{r_1}$ from $D^{(t)}$;
 Let $h^{(t)} = \arg \min_{h \in \mathcal{H}} \sum_{i=1}^{r_1} \ell(h, z_i)$ be the empirical minimizer of ℓ ;
 Sample r_2 datapoints $z_{D, (t-1)r_2+1}, \dots, z_{D, tr_2}$ from each $D \in \mathcal{D}$;
 Use the Hedge algorithm to get the next iterate $D^{(t+1)} \in \Delta(\mathcal{D})$, using learning rate α and observing the payoff $\tilde{\rho}^{(t)} : \mathcal{D} \rightarrow [0, 1]$ where $\tilde{\rho}^{(t)}(D) = \frac{1}{r_2} \sum_{i=(t-1)r_2+1}^{tr_2} \ell(h^{(t)}, z_{D,i})$;
end for
 Return $\bar{h}$: a uniform distribution over $h^{(1:T)}$;

A.1. Proof of Theorem 7 (Row 2 of Table 1)

Theorem 7 For any $\varepsilon, \alpha \in (0, 0.5)$, $\delta > 0$, $k \in \mathbb{Z}_+$ and binary class $\mathcal{H}$, the sample complexity of MDL, $m_{\mathcal{H}}(\varepsilon + \alpha \cdot \text{OPT}, \delta, k)$, is $\tilde{O}(\varepsilon^{-2} (k \log(k/\delta) + \alpha^{-2} \log(k)(\log(1/\varepsilon\delta) + \text{VC}(\mathcal{H})))$.

Proof Let d denote the VC dimension of $\mathcal{H}$. Without loss of generality, assume $\varepsilon \leq \alpha$. Consider Algorithm 1, fixing $T = \frac{\ln(k)}{\varepsilon\alpha}$, $r_1 = C_1 \frac{d + \ln(T/\delta)}{\varepsilon\alpha}$, and $r_2 = \left\lceil C_2 \frac{\ln(k/\delta)}{T\varepsilon^2} \right\rceil$.

Fact 8 The regret of the “adversary” in the game dynamics induced by Algorithm 1 satisfies

$$\text{Reg}_+(D^{(1:T)}, \{\mathcal{L}_{(\cdot)}(h^{(t)})\}^{(1:T)}) \leq \frac{\ln(k)}{\alpha} + T\varepsilon + \alpha \max_{D^* \in \mathcal{D}} \sum_{t=1}^T \mathcal{L}_{D^*}(h^{(t)}),$$

with probability at least $1 - 2\delta$ for some choice of universal constant C_2 .

Proof The mixture distributions $D^{(1:T)}$ result from applying Hedge to the payoff functions $\tilde{\rho}^{(1:T)}$. Hence, by Lemma 6,

$$\text{Reg}_+(D^{(1:T)}, \tilde{\rho}^{(1:T)}) \leq \frac{\ln(k)}{\alpha} + \alpha \max_{D^* \in \mathcal{D}} \sum_{t=1}^T \tilde{\rho}^{(t)}(D^*).$$

To prove generalization, we will break each timestep t into r_2 sub-timesteps. For every $j \in [Tr_2]$, we let $\tilde{D}^{(j)} = D^{(\lceil j/r_2 \rceil)}$ and define $\tilde{c}^{(j)}$ to be the cost function $D \mapsto \frac{1}{r_2} (1 - \ell(h^{(t)}, z_{D,j}))$. We can rewrite the adversary’s regret as $\text{Reg}_+(D^{(1:T)}, \tilde{\rho}^{(1:T)}) = \text{Reg}(\tilde{D}^{(1:Tr_2)}, \tilde{c}^{(1:Tr_2)})$. Further observe that, since $\mathbb{E}_{z_{D,j}} [\tilde{c}^{(j)}] = \mathcal{L}_D(h^{(\lceil j/r_2 \rceil)})$ for every $j \in [Tr_2]$ and $D \in \mathcal{D}$, the empirical regret is unbiased: $\text{Reg}_+(D^{(1:T)}, \{\mathcal{L}_{(\cdot)}(h^{(t)})\}^{(1:T)}) = \text{Reg}(\tilde{D}^{(1:Tr_2)}, \{\mathbb{E}_{z_{(\cdot),j}} [\tilde{c}^{(j)}]\}^{(1:Tr_2)})$. By Lemma 4,

$$\begin{aligned} & \left| \text{Reg}_+(D^{(1:T)}, \{\mathcal{L}_{(\cdot)}(h^{(t)})\}^{(1:T)}) - \text{Reg}_+(D^{(1:T)}, \tilde{\rho}^{(1:T)}) \right| \\ &= \left| \text{Reg}(\tilde{D}^{(1:Tr_2)}, \tilde{c}^{(1:Tr_2)}) - \text{Reg}(\tilde{D}^{(1:Tr_2)}, \{\mathbb{E}_{z_{(\cdot),j}} [\tilde{c}^{(j)}]\}^{(1:Tr_2)}) \right| \\ &\leq O\left(\sqrt{\ln(k/\delta)T/r_2}\right) = O(T\varepsilon/C_2), \end{aligned}$$

Algorithm 2 Personalized Algorithm.

Input: Hypotheses $\mathcal{H}$, distributions $\mathcal{D}$;

Initialize $\mathcal{D}^{(1)} = \mathcal{D}$;

for $t = 1, 2, \dots, \lceil \log(k) \rceil$ **do**

 Run Algorithm 3 on $\mathcal{D}^{(t)}, \mathcal{H}$ to obtain $h^{(t)}$;

 Sample $O(\varepsilon^{-2}(\ln(k \ln(k)/\delta)))$ datapoints $\mathbf{X}_t^D$ from each $D \in \mathcal{D}$;

 Let $\mathcal{D}^{(t+1)}$ consist of D where $\widehat{\mathcal{L}}_{\mathbf{X}_t^D}(h^{(t)}) > \text{Median}(\widehat{\mathcal{L}}_{\mathbf{X}_t^D}(h^{(t)}))_{D \in \mathcal{D}}$;

end for

For each $D \in \mathcal{D}$, find t_D where $D \in \mathcal{D}^{(t_D)}$ but $D \notin \mathcal{D}^{(t_D+1)}$. Return $(D, h^{(t_D)})_{D \in \mathcal{D}}$.

with probability at least $1 - \delta$. Similarly, with probability at least $1 - \delta$, $\max_{D^* \in \mathcal{D}} \sum_{t=1}^T \widetilde{\rho}^{(t)}(D^*) \leq O(\frac{T\varepsilon}{C_2}) + \max_{D^* \in \mathcal{D}} \sum_{t=1}^T \mathcal{L}_{D^*}(h^{(t)})$. A union bound yields the claimed fact. $\blacksquare$

Next, we observe that, at each timestep t , $h^{(t)}$ is the empirical risk minimizer of ℓ on $C_1(d + \log(T/\delta))/\varepsilon\alpha$ samples from $D^{(t)}$. For sufficiently large C_1 , by Lemma 5, $\mathcal{L}_{D^{(t)}}(h^{(t)}) \leq \varepsilon + (1 + \alpha) \min_{h^* \in \mathcal{H}} \mathcal{L}_{D^{(t)}}(h^*)$ with probability at least $1 - \delta/T$. By union bound, $\sum_{t=1}^T \mathcal{L}_{D^{(t)}}(h^{(t)}) \leq T\varepsilon + T(1 + \alpha)\text{OPT}$ with probability at least $1 - \delta$. Putting together the regret bounds for the learner and adversary,

$$\begin{aligned} (1 - \alpha) \max_{D^* \in \mathcal{D}} \mathcal{L}_{D^*}(\bar{h}) - 2\varepsilon &= (1 - \alpha) \max_{D^* \in \mathcal{D}} \left(\frac{1}{T} \sum_{t=1}^T \mathcal{L}_{D^*}(h^{(t)}) \right) - 2\varepsilon \\ &\leq \frac{1}{T} \sum_{t=1}^T \mathcal{L}_{D^{(t)}}(h^{(t)}) \\ &\leq \varepsilon + (1 + \alpha) \min_{h^* \in \mathcal{H}} \frac{1}{T} \sum_{t=1}^T \mathcal{L}_{D^{(t)}}(h^*) \\ &\leq \varepsilon + (1 + \alpha)\text{OPT}. \end{aligned}$$

We can simplify $\max_{D^* \in \mathcal{D}} \mathcal{L}_{D^*}(\bar{h}) \leq \frac{1}{1-\alpha}(3\varepsilon + (1 + \alpha)\text{OPT}) \leq 6\varepsilon + (1 + 4\alpha)\text{OPT}$. Reparameterizing $\varepsilon \rightarrow \frac{1}{6}\varepsilon$ and $\alpha \rightarrow \frac{1}{4}\alpha$ yields the desired claim. Our sample complexity is $(r_1 + kr_2) \times T$ and thus $O\left(\frac{k \ln(k/\delta)}{\varepsilon^2} + \frac{(d + \log(T/\delta)) \ln(k/\delta)}{\varepsilon^2 \alpha^2}\right)$. $\blacksquare$

A.2. Proof of Theorem 9

Theorem 9 For any $\varepsilon, \delta > 0$, $k \in \mathbb{Z}$ and binary class $\mathcal{H}$, the sample complexity $m_{\mathcal{H}}(\varepsilon, \delta, k)$ of personalized multi-distribution learning is $\widetilde{O}(\varepsilon^{-2} \ln(k)(\text{VC}(\mathcal{H}) \ln(\text{VC}(\mathcal{H})k/\varepsilon) + k \ln(k/\delta)))$.

We now turn to proving this result.

Lemma 10 Consider the multi-distribution learning problem $(\mathcal{D}, \mathcal{H}, \ell)$. For any $h \in \Delta \mathcal{H}$, there exists a $\mathcal{D}' \subseteq \mathcal{D}$ where $|\mathcal{D}'| \geq |\mathcal{D}|/2$ and $\max_{D \in (\Delta \mathcal{D})_2} \mathcal{L}_D(h) \geq \max_{D \in \mathcal{D}'} \mathcal{L}_D(h)$.

Algorithm 3 Multi-Distribution Learning Algorithm (Mid).

Input: Hypotheses $\mathcal{H}$, distributions $\mathcal{D}$;

Take $\varepsilon^{-1}C(d \log(d/\varepsilon) + \log(1/\delta))$ samples $x_1, \dots, x_N$ from $\text{Uniform}(\mathcal{D})$ and obtain a covering $\mathcal{H}'$ of $\mathcal{H}$ by projection: for every $y \in \{[h(x_1), \dots, h(x_N)] \mid h \in \mathcal{H}\}$, include in $\mathcal{H}'$ an arbitrary choice of $h \in \mathcal{H}$ such that $[h(x_1), \dots, h(x_N)] = y$;

Initialize Hedge iterate $D^{(1)}$ on $(\Delta\mathcal{D})_2$, that is the set of 2-smooth distributions on $\mathcal{D}$;

Initialize Hedge iterate $h^{(1)}$ on the simplex $(\Delta\mathcal{H}')$;

for $t = 1, 2, \dots, T$ **do**

 Use the Hedge algorithm to get the next iterate $h^{(t+1)} = \text{Hedge}_{\Delta(\mathcal{H}')}(\{\hat{c}^{(\tau)}\}^{(1:t)})$, where $\hat{c}^{(t)}(h) = \ell(h, z)$ and $z \sim D^{(t)}$;

 Sample a $D' \sim \text{Uniform}(\mathcal{D})$ and a datapoint $z \sim D$;

 Run Hedge algorithm to get the next iterate $D^{(t+1)} = \text{Hedge}_{(\Delta\mathcal{D})_2}(\{\hat{c}^{(\tau)}\}^{(1:t)})$, where $\hat{c}^{(t)}(D) = 1[D' = D] \cdot |\mathcal{D}| \cdot \Pr_{D^{(t)}}(D')(1 - \ell(h^{(t)}, z))$;

end for

Return a uniform distribution over $h^{(1:T)}$;

Proof Fix an $h \in \Delta\mathcal{H}$. Consider all strict minorities of $\mathcal{D}$: $\text{Min} := \{D' \subseteq \mathcal{D} \mid |D'| < |\mathcal{D}|/2\}$. Let $\mathcal{D}_{\text{MinHard}}$ denote the strict minority on which h does worst, and $\mathcal{D}_{\text{MinEasy}}$ denote the strict minority on which h does best:

$$\mathcal{D}_{\text{MinHard}} = \arg \max_{D^* \in \text{Min}} \frac{1}{|D^*|} \sum_{D \in D^*} \mathcal{L}(h), \quad \mathcal{D}_{\text{MinEasy}} = \arg \min_{D^* \in \text{Min}} \frac{1}{|D^*|} \sum_{D \in D^*} \mathcal{L}(h).$$

First, we observe that $\max_{D \in (\Delta\mathcal{D})_2} \mathcal{L}(h) \geq \mathbb{E}_{D \sim \text{Uniform}(\mathcal{D})} [\mathcal{L}(h) \mid D \notin \mathcal{D}_{\text{MinEasy}}]$, where $\text{Uniform}(\mathcal{D})$ is the uniform mixture over $\mathcal{D}$. Second, we observe that $\mathbb{E}_{D \sim \text{Uniform}(\mathcal{D})} [\mathcal{L}(h) \mid D \notin \mathcal{D}_{\text{MinEasy}}] \geq \max_{D \in \mathcal{D} \setminus \mathcal{D}_{\text{MinHard}}} \mathcal{L}_D(h)$. Thus, $D' = \mathcal{D} \setminus \mathcal{D}_{\text{MinHard}}$ satisfies the desired property. $\blacksquare$

Lemma 11 Consider a multi-distribution learning problem $(\mathcal{D}, \mathcal{H}, \ell)$. Algorithm 3 returns a hypothesis $\bar{h}$ such that with probability $1 - \delta$,

$$\max_{D^* \in (\Delta\mathcal{D})_2} \mathcal{L}_{D^*}(\bar{h}) \leq \min_{h^* \in \Delta\mathcal{H}} \max_{D^* \in (\Delta\mathcal{D})_2} \mathcal{L}_{D^*}(h^*) + \varepsilon.$$

It takes only $\tilde{O}(\varepsilon^{-2}(d \ln(dk/\varepsilon) + \ln(1/\delta)))$ samples.

Proof By construction, with probability at least $1 - \delta$, $\mathcal{H}'$ is an ε -net for $\mathcal{H}$ [HW86] under the distribution $\text{Uniform}(\mathcal{D})$. Consider any distribution $\mathcal{P} \in (\Delta\mathcal{D})_2$. Because any event that happens in $\mathcal{P}$ must also happen in $\text{Uniform}(\mathcal{D})$ with at least half the probability, including the event that $h(x) \neq h'(x)$, $\mathcal{H}'$ is a 2ε -net for $\mathcal{P}$. Since the range of ℓ is $[0, 1]$, it also follows that for any $h \in \mathcal{H}$, there is an $h' \in \mathcal{H}'$ such that $|\mathcal{L}_{\mathcal{P}}(h) - \mathcal{L}_{\mathcal{P}}(h')| < 2\varepsilon$. We now turn to arguing that our output $\text{Uniform}(h^{(1:T)})$ is nearly optimal for the discretized class $\mathcal{H}'$.

We observe that $D^{(1:T)}$ results from applying Hedge to (importance-weighted estimates of) stochastic cost functions that are bounded in $[0, 2]$. Moreover, the costs are bounded unbiased estimates of the true costs. Thus, as in our proof of Theorem 7, we can directly apply Hedge's regret bound (Lemma 3) and stochastic approximation (Lemma 4) to bound the adversary's regret

$\text{Reg}(D^{(1:T)}, \{1 - \mathcal{L}_{(\cdot)}(h^{(t)})\}^{(1:T)}) \leq O\left(\sqrt{\ln(k/\delta)T}\right)$. Note that this regret is defined only over the set $\Delta(\mathcal{D})_2$. Therefore choosing $T = \lceil C' \ln(k/\delta)/\varepsilon^2 \rceil$ for large C' gives $\text{Reg}(D^{(1:T)}, \{1 - \mathcal{L}_{(\cdot)}(h^{(t)})\}^{(1:T)}) \leq T\varepsilon$ with probability $1 - \delta$. Similarly, the learner's Hedge (Lemma 3) and the stochastic approximation (Lemma 4) gives the regret bound $\text{Reg}(h^{(1:T)}, \{\mathcal{L}_{D^{(t)}}(\cdot)\}^{(1:T)}) \leq O\left(\sqrt{\ln(|\mathcal{H}'|)T}\right)$. Note that this regret is defined only over the set $\mathcal{H}'$. Since $|\mathcal{H}'| \leq O((kN)^d)$ by Sauer Shelah's lemma, choosing $T \geq C\varepsilon^{-2}d \ln(dk \ln(d/\varepsilon\delta)/\varepsilon)$ for some large constant C guarantees that $\text{Reg}(h^{(1:T)}, \{\mathcal{L}_{D^{(t)}}(\cdot)\}^{(1:T)}) \leq T\varepsilon$ with probability at least $1 - \delta$. Putting together the regret bounds for the learner and adversary as before,

$$\begin{aligned}
 \max_{D^* \in \Delta(\mathcal{D})_2} \mathcal{L}_{D^*}(\bar{h}) - \varepsilon &\leq \frac{1}{T} \sum_{t=1}^T \mathcal{L}_{D^{(t)}}(h^{(t)}) \leq \varepsilon + \min_{h^* \in \mathcal{H}'} \frac{1}{T} \sum_{t=1}^T \mathcal{L}_{D^{(t)}}(h^*) \\
 &\leq 3\varepsilon + \min_{h^* \in \mathcal{H}} \frac{1}{T} \sum_{t=1}^T \mathcal{L}_{D^{(t)}}(h^*) \\
 &\leq 3\varepsilon + \text{OPT}.
 \end{aligned}$$

Our sample complexity is $2 \times T$ and thus $\tilde{O}(\varepsilon^{-2}(d \ln(dk/\varepsilon) + \ln(1/\delta)))$. ■

Proof [Proof of Theorem 9] Consider Algorithm 2. Let D^* be the product distribution of every $D \in \mathcal{D}$. Let d denote the VC dimension of $\mathcal{H}$. By Lemma 11, with probability at least $1 - \log(k)\delta$, for all $t \in [T]$,

$$\max_{D^* \in (\Delta\mathcal{D}^{(t)})_2} \mathcal{L}_{D^*}(h^{(t)}) \leq \min_{h^* \in \Delta\mathcal{H}} \max_{D^* \in (\Delta\mathcal{D}^{(t)})_2} \mathcal{L}_{D^*}(h^*) + \varepsilon.$$

By Lemma 10, there exists $\mathcal{D}' \subseteq \mathcal{D}^{(t)}$ where $|\mathcal{D}'| > |\mathcal{D}^{(t)}|/2$ and $\max_{D \in (\Delta\mathcal{D})_{\mathcal{P}^{(t)}}} \mathcal{L}_D(h^{(t)}) \geq \max_{D \in \mathcal{D}'} \mathcal{L}_D(h^{(t)})$. In other words, $\text{OPT} + \varepsilon \geq \max_{D \in \mathcal{D}'} \mathcal{L}_D(h^{(t)})$. By uniform convergence, with probability at least $1 - \delta$, for all $t \in [T]$ and $D \in \mathcal{D}^{(t)}$, $|\hat{\mathcal{L}}_D^{(t)}(h^{(t)}) - \mathcal{L}_D(h^{(t)})| \leq \varepsilon$. Thus, for every $D \in \mathcal{D}^{(t)} \setminus \mathcal{D}^{(t+1)}$, $\mathcal{L}_D(h^{(t)}) \leq 2\varepsilon + \text{OPT}$. Since the size of $\mathcal{D}^{(t)}$ is reduced by at least half every iteration, the algorithm terminates after $\lceil \ln(k) \rceil$ iterations. The algorithm's sample complexity comes from the samples needed for Lemma 10 and for evaluating each $h^{(t)}$, and taking a union bound over all iterations. ■