

THE COVERAGE PRINCIPLE: HOW PRE-TRAINING ENABLES POST-TRAINING

Anonymous authors

Paper under double-blind review

ABSTRACT

Language models demonstrate remarkable abilities when pre-trained on large text corpora and fine-tuned for specific tasks, but how and why pre-training shapes the success of the final model remains poorly understood. Notably, although pre-training success is often quantified by cross entropy loss, cross entropy can be poorly predictive of downstream performance. Instead, we provide a theoretical perspective on this relationship through the lens of *coverage*, which quantifies the probability mass the pre-trained model places on high-quality responses and which is necessary and sufficient for post-training and test-time scaling methods like Best-of-N to succeed. Our main results develop an understanding of *the coverage principle*, a phenomenon whereby next-token prediction implicitly optimizes toward a model with good coverage. In particular, we uncover a mechanism that explains the power of coverage in predicting downstream performance: *coverage generalizes faster than cross entropy*, avoiding spurious dependence on problem dependent parameters such as the sequence length. We also study practical algorithmic interventions with provable benefits for improving coverage, including (i) model/checkpoint selection procedures, (ii) gradient normalization schemes, and (iii) test-time decoding strategies.

1 INTRODUCTION

The remarkable capabilities of language models stem from a two-stage training process: (1) large-scale pre-training via next-token prediction with the cross-entropy loss (predicting what token should follow a prefix) and (2) targeted post-training—typically via reinforcement learning—to adapt the model to specific domains and tasks. Investing more compute and data into pre-training often enables post-training to produce a stronger model, but theoretical understanding of how these stages interact is limited. Indeed, despite substantial investment into scaling pre-training (Gadre et al., 2025; Sardana et al., 2024; Hoffmann et al., 2022), several works have demonstrated that starting post-training from a better next-token predictor does not ensure stronger performance on downstream tasks (Liu et al., 2022; Zeng et al., 2025; Chen et al., 2025; Lourie et al., 2025). Motivated by this disconnect, we theoretically investigate the connection between pre-training objectives and downstream success, asking:

Can we precisely characterize the relationship between the next-token prediction loss and downstream performance? What metrics are most predictive of downstream success?

Motivated by the recent interest in test-time scaling, we focus our attention on post-training via Best-of-N (BoN) sampling or reinforcement learning with verifiable rewards. For a prompt x , Best-of-N draws n responses y from the model and returns the best response according to a task-specific reward. Several prior works have demonstrated that the performance of BoN is strongly indicative of how well the model will perform after post-training via reinforcement learning (Yue et al., 2025; Wu et al., 2025).

Our starting point is the observation that cross-entropy alone cannot provide meaningful answers to the questions above; see Figure 1, which illustrates that cross-entropy can be *anti-correlated* with BoN performance, echoing Chen et al. (2025). Instead, we show that the missing link is the *coverage profile*, a novel refinement of cross-entropy that explicitly quantifies the model’s ability to assign sufficient probability mass to rare but high-quality responses.

Definition 1.1 (Coverage profile). *The coverage profile of a model $\hat{\pi}$ for a distribution π is*

$$\text{Cov}_N(\pi \parallel \hat{\pi}) := \mathbb{P}_{x \sim \mu, y \sim \pi(\cdot|x)} \left[\frac{\pi(y|x)}{\hat{\pi}(y|x)} \geq N \right], \quad (1)$$

where $N \geq 1$ is the number of Best-of-N sampling attempts.

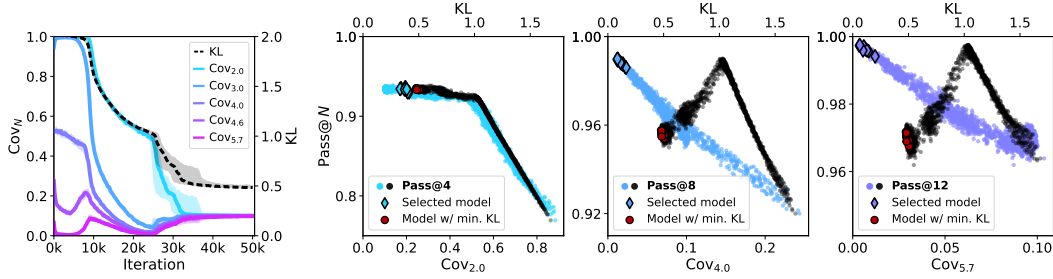


Figure 1: **The coverage profile predicts Pass@N better than KL divergence.** We train models in a graph reasoning task and record KL divergence, coverage profile (both measured w.r.t., π_D), and Pass@N performance; see Appendix G for details. Left: Convergence of coverage and KL divergence over training, showing that KL improves monotonically but coverage can *degrade* with training. Right: Scatter plots of KL (top axis), $\text{Cov}_{N/2}$ (lower axis) and Pass@N of checkpoints. Although KL and Cov_N exhibit comparable predictive power for small N , Cov_N is a better predictor for large N . Also visualized are checkpoints selected via the tournament procedure of Eq. (13) (marked $\diamond$) and by minimizing KL (marked red), demonstrating that the former selects better models for Pass@N.

Here, y is the full response when prompted with x , π represents the pre-training data distribution, which we presuppose covers downstream tasks of interest, and $\hat{\pi}$ is the pre-trained model. We prove that a **good coverage profile is necessary and sufficient for Best-of-N to succeed** (see Section 2, as well as Propositions D.6 and D.7). This is highlighted in Figure 1, where we find that the coverage profile is correlated with downstream performance for Best-of-N (which is exactly Pass@N), even when cross-entropy is not.¹ Motivated by this characterization of BoN performance, we ask: *When, and through what mechanism, does next-token prediction produce a model $\hat{\pi}$ with good coverage?*

1.1 CONTRIBUTIONS

We develop a theoretical understanding of **the coverage principle**, whereby next-token prediction implicitly optimizes toward a model with a good coverage profile, inheriting the training corpus’ coverage over tasks of interest.

Cross-entropy: Scaling laws and limitations (Section 3). We begin by deriving provable scaling laws that link cross-entropy—specifically, a certain sequence-level notion—to coverage and hence downstream performance, but show that cross-entropy can be sensitive to sequence length and other problem parameters, leading to vacuous predictions; this motivates our main results.

Next-token prediction implicitly optimizes coverage (Section 4). The first of our main theoretical results (Theorem 4.1) is a new generalization analysis for next-token prediction (more generally, maximum likelihood) that exploits the unique structure of the logarithmic loss to show that **coverage can generalize faster than cross entropy**; we refer to this as the coverage principle. Concretely, our analysis shows that the coverage profile for models learned with next-token prediction (i) avoids spurious dependence on problem-dependent parameters such as sequence length (in contrast to cross-entropy), and (ii) converges *faster* still as the tail parameter N is increased. Our analysis—which is similar in spirit to Mendelson’s *small ball method* (Mendelson, 2014; 2017)—can be viewed as a giving a new, fine-grained understanding of maximum likelihood.

Stochastic gradient descent through the lens of coverage (Section 5). The preceding results apply to general model classes II, but consider the empirical maximizer of the next-token prediction (maximum likelihood) objective, in the vein of classical techniques in learning theory. For the second of our main results, we focus on a specific model class—overparameterized autoregressive linear models (2)—but take a more realistic approach and analyze stochastic gradient descent (SGD) on the next-token prediction objective, in the one-pass (“compute-optimal”) regime. We show that while SGD provably optimizes the coverage profile, it experiences suboptimal dependence on the sequence length H . We then show that *gradient normalization* (which is loosely connected to Adam-like updates (Bernstein & Newhouse, 2024)) provably improves coverage, removing dependence on the sequence length.

¹Formally, the coverage profile refines cross-entropy/KL divergence; the former is the cumulative distribution function (CDF) of the log density ratio $\log \frac{\pi(y|x)}{\hat{\pi}(y|x)}$, while KL divergence is the mean; see Remark C.1.

Interventions for better coverage (Section 6). Finally, we look beyond standard next-token prediction and explore families of new interventions aimed at improving coverage in theory.

(i) Test-time (Section 6.1). We show that for standard token-level SGD, a new test-time decoding strategy inspired by *test-time training* (Sun et al., 2020; Akyürek et al., 2025) provably improves coverage.

(ii) Model/checkpoint selection (Section 6.2). For selecting the best model (or checkpoint) from a small number of candidates, we give *tournament* procedures that enjoy significantly better coverage profile (particularly with respect to the tail parameter N) than naïve validation with cross-entropy.

Additional results (Appendix F). Beyond the results above, we show that: (1) MLE can find models with low coverage even in the presence of severe misspecification; (2) coverage can generalize better under additional structural properties of the model class such as convexity (Appendix F.2).

In summary, we believe that coverage offers a new perspective on the connection between pre-training objectives and downstream post-training success. Our results demonstrate that this perspective is mathematically rich and fundamental, opening the door to a deeper understanding; cf. Appendix A.

2 PROBLEM SETUP

We now introduce the formal problem setup for the remainder of the paper.

Next-token prediction and maximum likelihood. We work in the following setting, which subsumes next-token prediction: $\mathcal{X}$ is the prompt space, $\mathcal{Y}$ is the response space, and $\pi_D : \mathcal{X} \rightarrow \Delta(\mathcal{Y})$ is the data distribution. We are given a dataset $\mathcal{D} = \{(x^i, y^i)\}_{i=1}^n$ where $x^i \sim \mu$ and $y^i \sim \pi_D(\cdot | x^i)$. We consider the maximum likelihood objective $\hat{L}_n(\pi) := \sum_{i=1}^n \log \pi(y^i | x^i)$, and refer to $\hat{\pi} := \arg \max_{\pi \in \Pi} \hat{L}_n(\pi)$ as the *maximum likelihood estimator* for a user-specified model class Π . This is a generalization of the next-token prediction, where $\mathcal{Y} = \mathcal{V}^H$ is a token sequence and $\pi(y | x) = \prod_{h=1}^H \pi(y_h | x, y_{1:h-1})$ is explicitly autoregressive, so that $\hat{L}_n(\pi) = \sum_{i=1}^n \sum_{h=1}^H \log \pi(y_h^i | x^i, y_{1:h-1}^i)$. We specialize to next-token prediction at certain points but otherwise focus on the general setting. We make the following realizability assumption throughout.

Assumption 2.1 (Realizability). *The data distribution π_D is realizable by some model $\pi \in \Pi$.*

This formulation captures pre-training and SFT, with some caveats; see Appendix A.1.

Post-training and the coverage profile. Given a reward function $r_T(x, y) \in \{0, 1\}$ representing success at a downstream task T , the goal is to fine-tune $\hat{\pi}$ —through reinforcement learning or test-time scaling—to obtain near-optimal reward. We show (Propositions D.6 and D.7) that for any task-specific comparator policy $\pi_T : \mathcal{X} \rightarrow \Delta(\mathcal{Y})$, Best-of-N sampling with $\tilde{\Theta}(N)$ samples satisfies $\mathbb{E}_{x \sim \mu} [r_T(x, \pi_T(x)) - r_T(x, \hat{\pi}_{N/10}^{\text{BoN}}(x))] \preceq \text{Cov}_N(\pi_T \| \hat{\pi})$, so a good coverage profile for π_T is sufficient for high reward. Moreover, in a worst-case sense a good coverage profile is necessary for high reward; see Proposition D.7. Further, while less well understood, some form of coverage is thought to be necessary for the success of reinforcement learning methods like GRPO (Yue et al., 2025; Song et al., 2024).

Returning to pre-training, it is clear that there is little hope that next-token prediction will produce a model $\hat{\pi}$ with good coverage with respect to a downstream task unless the data distribution π_D itself has reasonable coverage with respect to this task. We therefore posit that the data distribution covers such a downstream task, in the sense that it includes high-reward responses with some bounded-below probability. Since coverage satisfies a transitivity property, it follows that coverage with respect to π_D implies coverage with respect to the optimal policy for the downstream task. For example, if π_D has a 10% chance of generating a correct response, and $\text{Cov}_{N/10}(\pi_D \| \hat{\pi}) = \varepsilon$, then we get 10ε error.² Thus, **going forward, we focus on understanding when next-token prediction achieves good coverage $\text{Cov}_N(\pi_D \| \hat{\pi})$ relative to the data distribution π_D itself**, and avoid concerning ourselves with specific details of the task policy π_T or the specific relationship between π_T and π_D .

Autoregressive linear models. We analyze next-token prediction and maximum likelihood for general model classes Π , but our running example throughout the paper will be the class Π of *autoregressive linear models*, defined by a known feature map $\phi : \mathcal{X} \times \mathcal{V}^* \rightarrow \mathbb{R}^d$. For each parameter $\theta \in \Theta \subset \mathbb{R}^d$, the model $\pi_\theta = (\pi_\theta)_{h=1}^H$ is defined by

$$\pi_\theta(y_h | x, y_{1:h-1}) \propto \exp(\langle \theta, \phi(x, y_{1:h}) \rangle). \quad (2)$$

²See Proposition D.5 for formal results.

In practice, autoregressive sequence models—such as those based on transformers—generate each token by sampling from a softmax distribution whose logits are given by a linear combination of learned features (Radford et al., 2019). Eq. (2) simplifies this by freezing the feature map, yet remains expressive enough to model complex non-Markovian dependencies, depending on the choice of features.

Assumption 2.2. We assume that $\Theta \subseteq \{\theta : \|\theta\| \leq 1\}$ is convex, and $\sup_{h,x,y_{1:h}} \|\phi(x, y_{1:h})\| \leq B$.

3 CROSS-ENTROPY AND COVERAGE: SCALING LAWS AND LIMITATIONS

A natural approach to understanding when next-token prediction achieves good coverage is to appeal to cross-entropy—perhaps first showing that next-token prediction achieves low cross-entropy (which is true asymptotically), and then relating cross-entropy to coverage. In this section we motivate our main results by showing that while this is possible in a weak sense, it does not yield predictive guarantees for downstream performance in the finite-sample regime.

Define the *sequence-level* cross-entropy for $\hat{\pi}$ as $D_{\text{CE}}(\pi_D \parallel \hat{\pi}) := \mathbb{E}_{\pi_D} \left[\sum_{h=1}^H \log \frac{1}{\hat{\pi}(y_h | x, y_{1:h-1})} \right]$.

Since $\mathbb{E}_{\mathcal{D} \sim \pi_D} [\hat{L}_n(\pi)] = -n \cdot D_{\text{CE}}(\pi_D \parallel \pi)$, one expects that as we scale up compute, number of samples n , and model capacity Π , $D_{\text{CE}}(\pi_D \parallel \hat{\pi}) \rightarrow D_{\text{CE}}(\pi_D \parallel \pi_D)$, or equivalently $D_{\text{KL}}(\pi_D \parallel \hat{\pi}) \rightarrow 0$, where $D_{\text{KL}}(\pi_D \parallel \hat{\pi}) := \mathbb{E}_{\pi_D} \left[\sum_{h=1}^H \log \frac{\pi_D(y_h | x, y_{1:h-1})}{\hat{\pi}(y_h | x, y_{1:h-1})} \right]$ is the sequence-level KL divergence.

A simple scaling law for cross-entropy. We show below that if that the model $\hat{\pi}$ has reasonable KL divergence to the data distribution, the coverage profile can be bounded:

Proposition 3.1 (KL-to-coverage; see Proposition D.1). *For all $N \geq e$, $\text{Cov}_N(\pi_D \parallel \hat{\pi}) \leq \frac{D_{\text{KL}}(\pi_D \parallel \hat{\pi})}{\log(N/e)}$.*

Combining Proposition 3.1 with Proposition D.6 and our assumption that π_D has good coverage with respect to the downstream task yields a simple “scaling law” for test-time compute with BoN:

Consider a task of interest with reward $r_T(x, y)$, and suppose the data distribution π_D itself has constant probability of success (i.e., sampling $y \sim \pi_D(\cdot | x)$ with $r_T(x, y) = 1$). To achieve sub-optimality ε with Best-of-N, it suffices to choose the compute budget N as

$$N \approx \exp\left(\frac{D_{\text{KL}}(\pi_D \parallel \hat{\pi})}{\varepsilon}\right). \quad (3)$$

That is, for a fixed model $\hat{\pi}$ and KL-divergence level $D_{\text{KL}}(\pi_D \parallel \hat{\pi}) \leq D_{\text{CE}}(\pi_D \parallel \hat{\pi})$, Eq. (3) predicts that test-time compute should increase exponentially with the desired accuracy ε .³

Insufficiency of cross-entropy. At first glance, this seems to be in line with empirical test-time scaling laws (OpenAI, 2024), but there is an issue: While *token-level* cross-entropy has been observed to be modest in contemporary language models (Kaplan et al., 2020; Hoffmann et al., 2022; Xia et al., 2022), the *sequence-level* cross-entropy (and KL-divergence) generally grows with the length H of the sequence, so that Eq. (3) predicts exponential test-time scaling in the sequence length. Moreover, such a law cannot hold if we only assume token-level cross-entropy is bounded; see Proposition D.7.

Is this the end of the story? On the one hand, it is simple to show (Proposition D.2) that Proposition 3.1 is tight for a worst-case pair of models. Moreover, even for the autoregressive linear model in Eq. (2), sequence-level KL divergence scales linearly with the sequence length H , as shown in the next result.

Proposition 3.2. *Fix $H \in \mathbb{N}$. There exists $\phi : \mathcal{X} \times \mathcal{V}^* \rightarrow \mathbb{B}_2(1)$ and induced autoregressive linear class Π with parameter space $\Theta = \mathbb{B}_2(1)$, distribution μ over $\mathcal{X}$ and data distribution $\pi_D \in \Pi$, such that for any proper estimator $\hat{\pi} = \hat{\pi}(\mathcal{D}) \in \Pi$, it holds that w.p. at least 0.5, $D_{\text{KL}}(\pi_D \parallel \hat{\pi}) \geq \Omega(\frac{H}{n})$.*

This behavior is reflected empirically in Figure 2 in a graph reasoning task. Yet, for this task (Figure 2), we find that, in spite of large cross-entropy/KL, next-token prediction learns a model $\hat{\pi}$ with a good coverage profile across a range of sequence lengths and that downstream Best-of-N succeeds. Why is this happening? In light of the discussion above, it must be related to specific inductive bias of the next-token prediction objective itself.

³Neither KL divergence nor the coverage profile are observable quantities (though cross entropy is an estimable upper bound on KL), so this is a theoretical prediction rather than a practical one as-is; see Remark C.2.

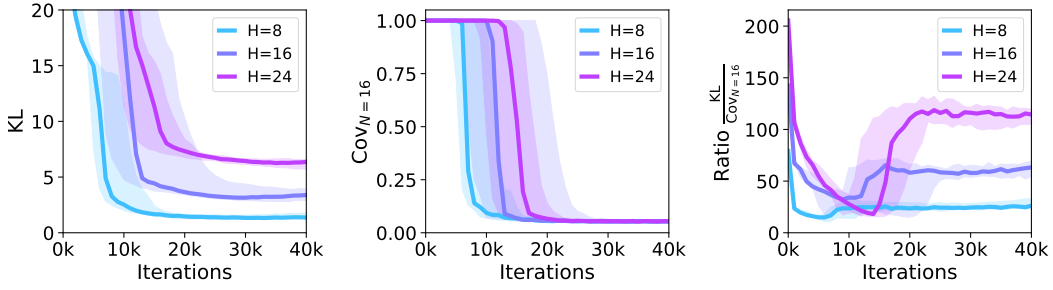


Figure 2: **The coverage profile avoids spurious dependence on sequence length.** We train models in a graph reasoning task and record their KL divergence and coverage profile, measured w.r.t., π_D as we vary the problem horizon (sequence length); see Appendix G for details. Left: Convergence of KL over training for three horizons H , demonstrating that KL at convergence scales linearly in the horizon H . Center: Convergence of Cov_N over training, manifesting no dependence on H at convergence. Right: Ratio of KL over Cov_N , showing that Proposition 3.1 can be overly conservative.

A glimmer of hope: Case study in Bernoulli models. To see why large cross-entropy may not be a barrier to coverage, consider perhaps the simplest setting, *Bernoulli models*, where $\mathcal{X} = \{\perp\}$, $\mathcal{Y} = \{0, 1\}$, $\Pi = \{\text{Ber}(p)\}_{p \in (0, 1/2)}$, and $\pi_D = \text{Ber}(p^*)$ for some small $p^* \in (0, 1/2)$.

The maximum likelihood model is $\hat{\pi} = \text{Ber}(\hat{p})$, where $\hat{p}$ is the empirical frequency of $y = 1$ in the dataset. We observe that with positive probability (and constant probability if $n \leq 1/p^*$), the dataset $\mathcal{D}$ will only contain examples where $y = 0$, so that the maximum likelihood model is $\hat{\pi} = \text{Ber}(0)$. This implies that expected KL divergence is infinite: $\mathbb{E}[D_{\text{KL}}(\pi_D \parallel \hat{\pi})] = +\infty$. However, the coverage profile turns out to be well-behaved; a direct calculation shows that $\text{Cov}_N(\pi_D \parallel \hat{\pi}) \lesssim \frac{\log(\delta^{-1})}{n}$ with probability at least $1 - \delta$ for all $N \geq 2$; this gives hope that even though cross-entropy itself is infinite, maximum likelihood may actually learn a model with good coverage in the background. In what follows, we will show that this is not a fluke, but a general phenomenon.

Remark 3.1 (Missing mass). *The underlying issue is one of missing mass: there are responses that even a well-generalizing learner will fail to cover, and for these we may incur a large contribution to the KL-divergence. More generally, KL-divergence and cross-entropy are susceptible to contributions of the scale $\log W_{\max}$ where $W_{\max} = \max_{\pi \in \Pi} \|\frac{\pi_D}{\pi}\|_{\infty}$ (which could be as large as H , as in Proposition 3.2) when the model does not have enough information to generalize/extrapolate. This phenomenon is particularly pronounced when the prompt distribution is heterogeneous.*

4 NEXT-TOKEN PREDICTION IMPLICITLY OPTIMIZES COVERAGE

We now present our main result, which establishes the *coverage principle*: due to the unique structure of the logarithmic loss, maximum likelihood can learn models with good coverage even when cross-entropy is vacuously large. We make use of the following covering number.

Definition 4.1. For a class Π and $\alpha \geq 0$, we let $\mathcal{N}_{\infty}(\Pi, \alpha)$ denote the size of the smallest cover Π' such that for all $\pi \in \Pi$, there exists $\pi' \in \Pi'$ such that $\sup_{x \in \mathcal{X}, y \in \mathcal{Y}} |\log \pi(y \mid x) - \log \pi'(y \mid x)| \leq \alpha$.

Theorem 4.1 (Coverage principle). Let $N \geq 8$ be given and $c > 0$ be an absolute constant. Suppose Assumption 2.1 holds. With probability at least $1 - \delta$, the maximum likelihood estimator satisfies

$$\text{Cov}_N(\hat{\pi}) \lesssim \frac{1}{\log N} \cdot \underbrace{\inf_{\varepsilon > 0} \left\{ \frac{\log \mathcal{N}_{\infty}(\Pi, \varepsilon) + \varepsilon}{n} \right\}}_{=: \mathcal{C}_{\text{fine}}(\Pi, n)} + \underbrace{\frac{\log \mathcal{N}_{\infty}(\Pi; c \log N) + \log(\delta^{-1})}{n}}_{=: \mathcal{C}_{\text{coarse}}(\Pi, N, n)}. \quad (4)$$

Theorem 4.1 has two terms: a *coarse-grained term* $\mathcal{C}_{\text{coarse}}(\Pi, N, n)$ and *fine-grained term* $\mathcal{C}_{\text{fine}}(\Pi, n)$; we interpret each term below.

Fine-grained term. $\mathcal{C}_{\text{fine}}(\Pi, n)$ evaluates the covering number $\mathcal{N}_{\infty}(\Pi, \varepsilon)$ at a small scale ε (typically $\varepsilon \approx \text{poly}(1/n)$), which matches typical bounds for conditional density estimation (e.g., Bilodeau et al. (2023)) in KL divergence; however, unlike KL-based bounds this term has *no explicit dependence on sequence length H or density ratios $\log W_{\max}$* . The term is further scaled by $1/\log N$, which implies that *coverage enjoys faster convergence as we move further into the tail by increasing N* ; this reflects the unique structure of the logarithmic loss, and may be viewed as a new form of implicit bias.

Summarizing, the fine-grained term witnesses the *coverage principle*: coverage enjoys faster generalization than cross-entropy; roughly, the rate is what we would expect (via [Proposition 3.1](#)) if we could somehow control KL without paying for the sequence length H or density ratio $\log W_{\max}$. See [Appendix E](#) for a detailed comparison to standard (asymptotic and non-asymptotic) generalization bounds for maximum likelihood based on Hellinger distance and KL-divergence.

Coarse-grained term. The coarse-grained term $\mathcal{C}_{\text{coarse}}(\Pi, N, n)$ captures the *missing mass* phenomenon exemplified by the Bernoulli example in the prequel. This term is not explicitly normalized by $1/\log N$ (compared to the fine-grained term), but depends on the covering number $\mathcal{N}_{\infty}(\Pi, \alpha)$ only at a very large scale $\alpha \approx \log N$. This implies that the dependence on the capacity of Π in this term vanishes as we increase N .

Overall, while the guarantee in [Eq. \(4\)](#) might look surprising at first glance (particularly the coarse term, as we are not aware of any existing generalization bounds with dependence on covering numbers at such a large scale), we show in [Proposition F.1](#) ([Appendix J](#)) that both terms are tight in general.

Overview of analysis. The proof of [Theorem 4.1](#) is given in [Appendix J](#) (with a high-level sketch in [Appendix J.1](#)). The basic idea is to interpret the condition $\text{Cov}_N(\pi) \geq \varepsilon$ as a small ball-like *anti-concentration* condition in the vein of [Mendelson \(2014; 2017\)](#). That is, for models π where coverage is large, the condition $\text{Cov}_N(\pi) \geq \varepsilon$ witnesses a *one-sided* bound which implies that the empirical likelihood of π is *not too large* with high probability, and thus π cannot be a maximum-likelihood solution.

The coarse-grained term $\mathcal{C}_{\text{coarse}}(\Pi, N, n)$ enters because we only need to show that the coverage profile concentrates, not the log loss itself. The fine-grained term $\mathcal{C}_{\text{fine}}(\Pi, n)$ enters from one-sided concentration of the empirical likelihood, with the $1/\log N$ scaling arising from the following form of implicit bias: If an example (x^i, y^i) is such that $\pi_0(y^i|x^i)/\pi(y^i|x^i) \geq N$, this witnesses a negative contribution of order $\log N$ to the difference $\hat{L}_n(\pi) - \hat{L}_n(\pi_0)$.

Discussion. We emphasize that while covering numbers are a fundamental and widely used notion of capacity in statistical learning and estimation ([van de Geer, 2000](#); [Zhang, 2002](#); [Rakhlin & Sridharan, 2012](#); [Bilodeau et al., 2023](#)), they are conservative from a modern generalization perspective. Nonetheless, [Theorem 4.1](#) shows that they are sufficient to capture rich aspects of generalization for coverage, and we expect that our core analysis techniques can be combined with contemporary advances in generalization theory for overparameterized models ([Belkin et al., 2019](#); [Bartlett et al., 2020](#)).

4.1 EXAMPLES

To build intuition, we analyze the behavior of [Theorem 4.1](#) under a growth assumption on the covering number, then discuss how autoregressive linear models exemplify the coverage principle.

Corollary 4.1. (i) Parametric regime: Suppose that there are parameters $d \geq 2$ and $C \geq 2$ such that $\log \mathcal{N}_{\infty}(\Pi, \alpha) \leq d \log(C/\alpha)$ for $\alpha \in (0, C/2]$. Then for any $N \geq 8$, with probability at least

$$1 - \delta, \text{Cov}_N(\hat{\pi}) \lesssim \frac{d[\log(C/\log N)]_+ + \frac{\log(Cn)}{\log N} + \log(1/\delta)}{n}.$$

(ii) Nonparametric regime: Suppose that there are parameters $C \geq 2$ and $p > 0$ such that $\log \mathcal{N}_{\infty}(\Pi, \alpha) \leq (C/\alpha)^p$ for $\alpha \in (0, C/2]$. Then for any $N \geq 8$ and $n \geq \log^{1/p} N \cdot (C/\log N)^p$, with probability at least $1 - \delta$, $\text{Cov}_N(\hat{\pi}) \lesssim \frac{1}{\log N} \left(\frac{C^p}{n}\right)^{\frac{1}{p+1}} + \frac{\log(1/\delta)}{n}$.

This result shows that for sufficiently rich classes (e.g., when $p > 0$), the fine-grained term dominates the coarse-grained term for n sufficiently large. On the other hand, for simple classes (e.g., when $p = 0$), the coarse-grained term can dominate the fine-grained term.

Autoregressive linear models: Low dimension. We now specialize to our running example, the autoregressive linear model in [Eq. \(2\)](#). This class satisfies $\log \mathcal{N}_{\infty}(\Pi, \alpha) \asymp d \log(BH/\alpha)$ (corresponding to the parametric regime in [Corollary 4.1](#)), and so, coverage generalizes in a (nearly) horizon-independent fashion, in stark contrast to the cross-entropy lower bound in [Proposition 3.2](#). The only drawback (which is fundamental) is that since the class has low capacity, the coarse-grained term dominates for most parameter regimes, and the improvement as N scales is quite modest.

Autoregressive linear models: High dimension. As a more interesting example, we next look at the behavior of next-token prediction for autoregressive linear models in an “overparameterized” regime where the dimension d is arbitrarily large ([Zhang, 2002](#); [Neyshabur et al., 2015](#); [Bartlett et al., 2017](#)); here we expect polynomial dependence on the norm parameter B , as it is the only parameter

that controls the richness of the class Π . In this regime, it turns out that in the worst-case, the capacity $\log \mathcal{N}_\infty(\Pi, \alpha)$ scales polynomially in H . To address, this we prove a refined version of [Theorem 4.1](#) that adapts to the variance in the data distribution π_D , avoiding explicit dependence on sequence length.

Define the *inherent variance* for the data distribution as

$$\sigma_\star^2 := \mathbb{E}_{\pi_D} \left[\sum_{h=1}^H \left\| \phi(x, y_{1:h}) - \bar{\phi}_{\pi_D}(x, y_{1:h-1}) \right\|^2 \right], \quad (5)$$

where $\bar{\phi}_{\pi_D}(x, y_{1:h-1}) := \mathbb{E}_{y_h \sim \pi_D(\cdot | x, y_{1:h-1})} [\phi(x, y_{1:h})]$ is the average feature vector given the prefix $(x, y_{1:h-1})$. We can interpret the inherent variance $\sigma_\star^2$ as a notion of effective sequence length; it captures the number tokens that are “pivotal” in the sense that they have high variation conditioned on the prefix; the name reflects the observed phenomenon that, in language modeling, most tokens are near-deterministic given their prefix, with only a few having high entropy ([Abdin et al., 2024](#)). Thus, while $\sigma_\star^2$ can be as large as $B^2 H$ in the worst case, we expect it to be smaller in general.

Theorem 4.2 (Overparameterized autoregressive linear models). *Consider the autoregressive linear model (2), and suppose [Assumptions 2.1](#) and [2.2](#) hold. For any $N \geq 2$, next-token prediction achieves*

$$\mathbb{E}[\text{Cov}_N(\hat{\pi})] \lesssim \sqrt{\frac{\sigma_\star^2}{n \cdot \log N}} + \frac{B^2}{n}. \quad (6)$$

Similar to [Theorem 4.1](#), the first term in [Eq. \(6\)](#) can be viewed as “fine-grained”, but decreases with the tail parameter N , while the second “coarse-grained” term does not decrease with N but will typically be smaller to begin with. We prove (details in [Proposition K.1](#)) that this result is tight in the sense that if $\sigma_\star^2 \asymp H$, $n \geq H$ is indeed necessary to achieve good coverage in the high-dimensional regime.

We view the introduction of the inherent variance $\sigma_\star^2$ as an instance-dependent notion of complexity for autoregressive models to be a non-trivial conceptual contribution, which may find broader use.

5 STOCHASTIC GRADIENT DESCENT THROUGH THE LENS OF COVERAGE

The coverage-based generalization guarantees for next-token prediction in the prequel apply to general model classes Π , but consider the empirical maximizer $\hat{\pi} = \arg \max_{\pi \in \Pi} \hat{L}_n(\pi)$ of the next-token prediction (maximum likelihood) objective, in the vein of classical techniques in learning theory. For our second set of main results, we focus on autoregressive linear models (2) but take a more realistic approach and analyze stochastic gradient descent (SGD) in the single-pass regime. This setup is motivated by contemporary (“compute-optimal”) language model training, which typically uses one or fewer passes over the training corpus ([Kaplan et al., 2020](#); [Hoffmann et al., 2022](#)).

5.1 STOCHASTIC GRADIENT DESCENT HAS SUBOPTIMAL COVERAGE

For the next-token prediction objective, single-pass stochastic gradient descent (SGD) takes the form⁴

$$\theta^{t+1} \leftarrow \text{Proj}_\Theta(\theta^t + \eta \nabla \log \pi_{\theta^t}(y^t | x^t)), \quad (7)$$

for $x^t \sim \mu$ and $y^t \sim \pi_D(\cdot | x^t)$, where $\eta > 0$ is the learning rate. As the next-token prediction loss $L(\theta) := \mathbb{E}_{\pi_D}[-\log \pi_\theta(y | x)]$ is convex under this parameterization, we can show that SGD converges to π_D in KL divergence. This implies a coverage bound, albeit a suboptimal one.

Proposition 5.1 (SGD for autoregressive linear models). ***Upper bound:** Suppose [Assumptions 2.1](#) and [2.2](#) hold. As long as $\eta \leq \frac{1}{16HB^2}$, it holds that $\mathbb{E}[\frac{1}{T} \sum_{t=1}^T D_{\text{KL}}(\pi_D \| \pi_{\theta^t})] \leq \frac{1}{\eta T} + 4\eta\sigma_\star^2$. Choosing η to minimize this bound gives*

$$\mathbb{E}[\frac{1}{T} \sum_{t=1}^T \text{Cov}_N(\pi_{\theta^t})] \lesssim \frac{1}{\log N} \cdot \left(\sqrt{\frac{\sigma_\star^2}{T}} + \frac{B^2 H}{T} \right). \quad (8)$$

Lower bound: Suppose that $B \geq c \cdot \log^2(TH)$. Then there exists an autoregressive linear class Π such that for any constant step size $\eta > 0$, there exists an instance $\pi_D \in \Pi$ with $\sigma_\star \leq 1$ such that with probability at least 0.5, the SGD iterates satisfy $\text{Cov}_N(\pi_D \| \pi_{\theta^t}) \geq c \cdot \min\left\{\frac{H}{T \log N}, 1\right\}$ for any $t \in [T]$.

The coverage bound in [Eq. \(8\)](#) (which follows by passing from KL to coverage through [Proposition 3.1](#)) is similar to [Theorem 4.2](#), except that the second term $\frac{B^2 H}{T}$ has an unfortunate dependence on the sequence length H . The lower bound shows that this dependence is tight, and SGD can indeed

⁴ $\text{Proj}_\Theta(\cdot)$ denotes Euclidean projection onto Θ , so this is SGD on the loss $L(\theta) := \mathbb{E}[-\log \pi_\theta(y | x)]$.

experience poor coverage. The failure of SGD in [Proposition 5.1](#) is related to *heterogeneity* across prompts: there are some prompts for which the effective scale of the gradient in [Eq. \(7\)](#) grows with H , leading to divergence unless we use a small learning rate $\eta \lesssim \frac{1}{HB}$. Yet for other prompts, the effective gradient range is small, leading to slow convergence (on the order of $\Omega(H)$ steps) unless $\eta \gg \frac{1}{HB}$.

Remark 5.1 (Sequence-level SGD). *The update in [Eq. \(7\)](#) can be interpreted as a “sequence-level” form of SGD, since we perform a single gradient step for each full sequence y^t (note that $\nabla \log \pi_{\theta^t}(y^t \mid x^t) = \sum_{h=1}^H \nabla \log \pi_{\theta^t}(y_h^t \mid x^t, y_{1:h-1}^t)$). We view this as a model for what is done in practice, whereby one performs SGD on sequences of tokens spanning some fixed context window. While this context window may be shorter than the full training example (e.g., a long article), understanding the implications of a limited context window is beyond the scope of this work.*

5.2 GRADIENT NORMALIZATION IMPROVES COVERAGE

To address the suboptimality of SGD, we consider *gradient normalization* as a simple intervention. For a mini-batch $\mathcal{D} = \{(x^i, y^i)\}_{i=1}^K$ of K samples from $\pi_{\mathcal{D}}$, define the batch stochastic gradient as $\hat{g}(\theta; \mathcal{D}) = \frac{1}{|\mathcal{D}|} \sum_{(x,y) \in \mathcal{D}} \nabla \log \pi_{\theta}(y \mid x)$. We consider the following normalized SGD update:

$$\theta^{t+1} \leftarrow \text{Proj}_{\Theta} \left(\theta^t + \eta \cdot \frac{\hat{g}(\theta^t; \mathcal{D}^t)}{\lambda + \|\hat{g}(\theta^t; \mathcal{D}^t)\|} \right); \quad (9)$$

here $\mathcal{D}^t$ is a mini-batch with K fresh samples drawn i.i.d. from $\pi_{\mathcal{D}}$, and $\lambda > 0$ is a regularization parameter for numerical stability. We show that this update achieves a horizon-independent coverage bound.

Theorem 5.1. *Suppose [Assumption 2.1](#) and [Assumption 2.2](#) hold. Let $T, K \geq 1, N \geq 3$ be given. For an appropriate choice of $\eta, \lambda > 0$, the normalized SGD update (9) achieves the following bound:*

$$\mathbb{E} \left[\frac{1}{T} \sum_{t=1}^T \text{Cov}_N(\pi_{\theta^t}) \right] \lesssim \sqrt{\frac{\sigma_*^2}{T \cdot \log N}} + \frac{B^2}{T} + \frac{B}{K \cdot \log N}. \quad (10)$$

To achieve $\mathbb{E}[\text{Cov}_N(\hat{\pi})] \leq \varepsilon$ for a target level $\varepsilon > 0$, it suffices to choose $T = O\left(\frac{\sigma_*^2}{\varepsilon^2 \log N} + \frac{B^2}{\varepsilon}\right)$, $K = O\left(\frac{B}{\varepsilon \log N} + 1\right)$, giving total sample complexity $n = TK = O\left(\frac{\sigma_*^2 B}{\varepsilon^3 \log^2 N} + \frac{B^3 + \sigma_*^2}{\varepsilon^2 \log N} + \frac{B^2}{\varepsilon}\right)$.

[Theorem 5.1](#) shows that gradient normalization achieves horizon-independent coverage with a qualitatively similar rate to the guarantee for next-token prediction in [Theorem 4.2](#): To achieve coverage ε , both rates scale as $\text{poly}\left(\frac{\sigma_*^2}{\log N}, B, \varepsilon^{-1}\right)$, though the dependence on ε for [Theorem 5.1](#) is worse. We view this as another instance of the coverage principle, as the rate achieved by gradient normalization goes beyond what can be achieved by passing through KL divergence. We emphasize that minibatching alone is not enough to achieve this result; rather, minibatching is necessary to avoid excessive bias once we introduce gradient normalization.

As a remark, the normalized SG update in (9) is closely related to SignSGD ([Balles & Hennig, 2018](#)) and Adam ([Kingma & Ba, 2015](#)) as shown by [Bernstein & Newhouse \(2024\)](#). We believe that similar coverage guarantees could potentially be shown for these methods using our techniques.

Distillation. As an additional result, we show ([Theorem F.2](#) in [Appendix F.3](#)) that for a *distillation* setting, where $\pi_{\mathcal{D}}$ corresponds to a teacher model and we have access to its per-token logits, we can derive an improved gradient normalization scheme that fully closes the gap with [Theorem 4.2](#).

6 INTERVENTIONS FOR BETTER COVERAGE

In this section, we develop new interventions that improve coverage (and downstream performance) beyond the conventional algorithms analyzed in [Sections 4](#) and [5](#). We view these results as promising proofs of concept, opening the door for further research into interventions driven by coverage.

6.1 IMPROVING COVERAGE AT TEST TIME

In this section, we show that a new test-time decoding strategy inspired by *test-time training* ([Sun et al., 2020](#); [Akyürek et al., 2025](#)) leads to improved coverage when combined with token-level SGD.

We begin by departing from [Eq. \(7\)](#) and learning models with a *token-level* SGD update, defined as

$$\theta^{t,h+1} = \text{Proj}_{\Theta} \left(\theta^{t,h} + \eta \nabla \log \pi_{\theta^{t,h}}(y_h^t \mid x^t, y_{1:h-1}^t) \right), \text{ for } h = 0, \dots, H-1, \quad (11)$$

and $\theta^{t+1} \equiv \theta^{t+1,0} := \theta^{t,H}$ for $t \in [T]$, and where $(x^t, y_{1:H}^t) \sim \pi_D$. Below we show that, when combined with a test-time training-like update that performs token-level gradient updates *during test time*, the updates in Eq. (11) can circumvent the H -dependence in the lower bound of Proposition 5.1.

Concretely, we consider a distribution $\pi_\theta^{\text{TTT}} : \mathcal{X} \rightarrow \Delta(\mathcal{Y}^H)$ formally introduced in Appendix L.1, which can be interpreted as an augmented version of the autoregressive linear model π_θ that uses test-time training to sample. Given a prompt x , we first sample $y_1 \sim \pi_\theta(\cdot | x)$, then perform a gradient step $\theta' \leftarrow \text{Proj}_\Theta(\theta + \eta \nabla \log \pi_\theta(y_1 | x))$ to increase the probability of the token we just sampled. We then sample $y_2 \sim \pi_{\theta'}(\cdot | x, y_1)$, update $\theta'' \leftarrow \text{Proj}_\Theta(\theta' + \eta \nabla \log \pi_{\theta'}(y_2 | x, y_1))$, and so on. Once the full sequence $y_{1:H}$ is sampled, we reset back to θ , to process the next example at test-time. We show that when augmented with this test-time sampling scheme, token-level SGD achieves a horizon-independent coverage bound that matches and even slightly improves upon the bound for next-token prediction in Theorem 4.2.

Theorem 6.1 (Token-level SGD with test-time training). *Suppose Assumption 2.1 and Assumption 2.2 hold. For a suitably chosen parameter $\eta > 0$, token-level SGD (11) achieves*

$$\mathbb{E}\left[\frac{1}{T} \sum_{t=1}^T D_{\text{KL}}(\pi_D \| \pi_{\theta^t}^{\text{TTT}})\right] \lesssim \sqrt{\frac{\sigma_D^2}{T} + \frac{B^2}{T}}, \text{ and thus } \mathbb{E}\left[\frac{1}{T} \sum_{t=1}^T \text{Cov}_N(\pi_{\theta^t}^{\text{TTT}})\right] \lesssim \frac{1}{\log N} \left(\sqrt{\frac{\sigma_D^2}{T} + \frac{B^2}{T}}\right).$$

This improves Theorem 4.2 by a factor of $1/\sqrt{\log N}$ on the leading term and a factor of $1/\log N$ on the second term. Furthermore, the algorithm bypasses the lower bound on KL divergence for proper methods (Proposition 3.2), demonstrating a provable benefit of being *improper*.

6.2 SELECTING FOR COVERAGE

We now consider the problem of selecting a model (e.g., checkpoint) from a small number of candidates to achieve the best coverage. We introduce a tournament-like procedure that improves upon maximum likelihood in that it removes the requirement that $\pi_D \in \Pi$; it is guaranteed to find a model in the class with good coverage if one exists, even if π_D itself is not in the class. As an algorithmic intervention, we envision using this procedure to select a single training checkpoint or hyperparameter configuration to use for RL fine-tuning or test-time scaling. Indeed, as demonstrated in Figure 1, using cross-entropy as a selection criterion—as is standard—may result in poor coverage, and these procedures can be used to select better checkpoints. Our results here concern the general setting in Section 2, and are not restricted to autoregressive linear models.

A simple tournament estimator for coverage. Given a dataset $\mathcal{D} = \{(x^i, y^i)\}_{i \in [n]}$, define

$$\widehat{\text{Cov}}_N(\pi' \| \pi) := \frac{1}{n} \left| \left\{ i \in [n] : \frac{\pi'(y^i | x^i)}{\pi(y^i | x^i)} \geq N \right\} \right|, \quad (12)$$

which can be interpreted as an empirical version of the coverage profile $\text{Cov}_N(\pi' \| \pi)$ in Eq. (1) when $\pi' = \pi_D$. For $N \geq 1$, we consider the estimator

$$\widehat{\pi} := \arg \min_{\pi \in \Pi} \max_{\pi' \in \Pi} \widehat{\text{Cov}}_N(\pi' \| \pi). \quad (13)$$

Intuitively, this estimator chooses the model π that minimizes the maximum coverage against any other model π' in the class Π . When Π is small, we can implement this tournament by simply evaluating the empirical coverage in Eq. (12) for each pair. The main guarantee for this estimator is as follows.

Theorem 6.2. *Let $N \geq 1$ be given. Then, for any $a \in [0, 1]$, with probability at least $1 - \delta$, the tournament estimator (13) achieves*

$$\text{Cov}_{N^{1+a}}(\widehat{\pi}) \lesssim \min_{\pi \in \Pi} \text{Cov}_{N^a}(\pi) + \frac{1}{N^{1-a}} + \frac{\log(|\Pi|/\delta)}{n}. \quad (14)$$

This shows that the tournament achieves a coverage profile nearly as good as the best-in class, except for a small polynomial blow up, in that we bound the coverage at level N^{1+a} in terms of the coverage for the best-in class at level N^a .

Infinite class and improving the tournament. Eq. (13) can also be applied to general, infinite classes Π . In this case, it turns out that it improves upon the coverage achieved by the maximum likelihood estimator in Theorem 4.1 (see Theorem 6.2'). Furthermore, in Appendix F.4, we describe an improved tournament estimator that is able to remove the $1/N^{1-a}$ term from Theorem 6.2, thereby achieving nontrivial guarantees even when the coverage parameter N is a constant.

DISCUSSION AND FUTURE WORK

See Appendix A for discussion and open problems, and Appendix F for additional results.

REPRODUCIBILITY STATEMENT

We provide full proofs for all theoretical results in the appendix. [Appendix G](#) includes extensive experiment setup and implementation details for all empirical results. The source code is included in the supplementary material, along with the plotting scripts and data to reproduce [Figure 1](#) and [Figure 2](#).

REFERENCES

- Marah Abdin, Jyoti Aneja, Harkirat Behl, Sébastien Bubeck, Ronen Eldan, Suriya Gunasekar, Michael Harrison, Russell J Hewett, Mojan Javaheripi, Piero Kauffmann, et al. Phi-4 technical report. *arXiv preprint arXiv:2412.08905*, 2024.
- Ekin Akyürek, Mehul Damani, Adam Zweiger, Linlu Qiu, Han Guo, Jyothish Pari, Yoon Kim, and Jacob Andreas. The surprising effectiveness of test-time training for few-shot learning. In *Forty-second International Conference on Machine Learning*, 2025.
- Zeyuan Allen-Zhu and Yuanzhi Li. Physics of language models: Part 3.3, knowledge capacity scaling laws. In *The Thirteenth International Conference on Learning Representations*, 2025. URL <https://openreview.net/forum?id=FxNNiUgtfa>.
- Gregor Bachmann and Vaishnavh Nagarajan. The pitfalls of next-token prediction. *arXiv preprint arXiv:2403.06963*, 2024.
- Lukas Balles and Philipp Hennig. Dissecting adam: The sign, magnitude and variance of stochastic gradients. In *International Conference on Machine Learning*, pp. 404–413. PMLR, 2018.
- Hritik Bansal, Arian Hosseini, Rishabh Agarwal, Vinh Q. Tran, and Mehran Kazemi. Smaller, weaker, yet better: Training LLM reasoners via compute-optimal sampling. In *The Thirteenth International Conference on Learning Representations*, 2025. URL <https://openreview.net/forum?id=30yaXFQuD1>.
- Peter L. Bartlett and Andrea Montanari. Deep learning: A statistical viewpoint. *Acta Numerica*, 30: 87–201, 2021.
- Peter L. Bartlett, Dylan J. Foster, and Matus J. Telgarsky. Spectrally-normalized margin bounds for neural networks. In *Advances in Neural Information Processing Systems*, 2017.
- Peter L. Bartlett, Philip M. Long, Gábor Lugosi, and Alexander Tsigler. Benign overfitting in linear regression. *Proceedings of the National Academy of Sciences (PNAS)*, 117(48):30063–30070, 2020.
- Mikhail Belkin, Daniel Hsu, Siyuan Ma, and Soumik Mandal. Reconciling modern machine-learning practice and the classical bias–variance trade-off. *Proceedings of the National Academy of Sciences (PNAS)*, 116(32):15849–15854, 2019.
- Jeremy Bernstein and Laker Newhouse. Old optimizer, new norm: An anthology. In *OPT 2024: Optimization for Machine Learning*, 2024.
- Blair Bilodeau, Dylan J Foster, and Daniel M Roy. Minimax rates for conditional density estimation via empirical entropy. *The Annals of Statistics*, 2023.
- Adam Block and Yury Polyanskiy. The sample complexity of approximate rejection sampling with applications to smoothed online learning. In *The Thirty Sixth Annual Conference on Learning Theory*, pp. 228–273. PMLR, 2023.
- Bradley Brown, Jordan Juravsky, Ryan Saul Ehrlich, Ronald Clark, Quoc V Le, Christopher Re, and Azalia Mirhoseini. Large language monkeys: Scaling inference compute with repeated sampling, 2025. URL <https://openreview.net/forum?id=0xUEBQV54B>.
- Feng Chen, Allan Raventos, Nan Cheng, Surya Ganguli, and Shaul Druckmann. Rethinking fine-tuning when scaling test-time compute: Limiting confidence improves mathematical reasoning. *arXiv preprint arXiv:2502.07154*, 2025.

- Jinglin Chen and Nan Jiang. Information-theoretic considerations in batch reinforcement learning. In *International conference on machine learning*, pp. 1042–1051. PMLR, 2019.
- Yangyi Chen, Binxuan Huang, Yifan Gao, Zhengyang Wang, Jingfeng Yang, and Heng Ji. Scaling laws for predicting downstream performance in llms. *Transactions on Machine Learning Research*, 2024.
- Tianzhe Chu, Yuexiang Zhai, Jihan Yang, Shengbang Tong, Saining Xie, Dale Schuurmans, Quoc V Le, Sergey Levine, and Yi Ma. SFT memorizes, RL generalizes: A comparative study of foundation model post-training. In *Forty-second International Conference on Machine Learning*, 2025. URL <https://openreview.net/forum?id=dYur3yabMj>.
- Rick Durrett. *Probability: theory and examples*, volume 49. Cambridge university press, 2019.
- Amir-massoud Farahmand, Csaba Szepesvári, and Rémi Munos. Error propagation for approximate policy and value iteration. *Advances in Neural Information Processing Systems*, 2010.
- Bahare Fatemi, Jonathan Halcrow, and Bryan Perozzi. Talk like a graph: Encoding graphs for large language models. In *The Twelfth International Conference on Learning Representations*, 2024.
- Marc Finzi, Sanyam Kapoor, Diego Granziol, Anming Gu, Christopher De Sa, J Zico Kolter, and Andrew Gordon Wilson. Compute-optimal llms provably generalize better with scale. *arXiv preprint arXiv:2504.15208*, 2025.
- Dylan J Foster, Sham M Kakade, Jian Qian, and Alexander Rakhlin. The statistical complexity of interactive decision making. *arXiv:2112.13487*, 2021.
- Dylan J Foster, Akshay Krishnamurthy, David Simchi-Levi, and Yunzong Xu. Offline reinforcement learning: Fundamental barriers for value function approximation. In *Conference on Learning Theory*, pp. 3489–3489. PMLR, 2022.
- Dylan J Foster, Adam Block, and Dipendra Misra. Is behavior cloning all you need? understanding horizon in imitation learning. *arXiv preprint arXiv:2407.15007*, 2024.
- Dylan J Foster, Zakaria Mhammedi, and Dhruv Rohatgi. Is a good foundation necessary for efficient reinforcement learning? the computational role of the base model in exploration. *Conference on Learning Theory (COLT)*, 2025.
- Samir Yitzhak Gadre, Georgios Smyrnis, Vaishal Shankar, Suchin Gururangan, Mitchell Wortsman, Rulin Shao, Jean Mercat, Alex Fang, Jeffrey Li, Sedrick Keh, et al. Language models scale reliably with over-training and on downstream tasks. In *The Thirteenth International Conference on Learning Representations*, 2024.
- Samir Yitzhak Gadre, Georgios Smyrnis, Vaishal Shankar, Suchin Gururangan, Mitchell Wortsman, Rulin Shao, Jean Mercat, Alex Fang, Jeffrey Li, Sedrick Keh, Rui Xin, Marianna Nezhurina, Igor Vasiljevic, Luca Soldaini, Jenia Jitsev, Alex Dimakis, Gabriel Ilharco, Pang Wei Koh, Shuran Song, Thomas Kollar, Yair Carmon, Achal Dave, Reinhard Heckel, Niklas Muennighoff, and Ludwig Schmidt. Language models scale reliably with over-training and on downstream tasks. In *The Thirteenth International Conference on Learning Representations*, 2025. URL <https://openreview.net/forum?id=iZeQBqJamf>.
- Kanishk Gandhi, Ayush Chakravarthy, Anikait Singh, Nathan Lile, and Noah D Goodman. Cognitive behaviors that enable self-improving reasoners, or, four habits of highly effective stars. *arXiv preprint arXiv:2503.01307*, 2025.
- Zhaolin Gao, Jonathan D Chang, Wenhao Zhan, Owen Oertell, Gokul Swamy, Kianté Brantley, Thorsten Joachims, J Andrew Bagnell, Jason D Lee, and Wen Sun. REBEL: Reinforcement learning via regressing relative rewards. *arXiv:2404.16767*, 2024.
- Behrooz Ghorbani, Orhan Firat, Markus Freitag, Ankur Bapna, Maxim Krikun, Xavier Garcia, Ciprian Chelba, and Colin Cherry. Scaling laws for neural machine translation. In *International Conference on Learning Representations*, 2022. URL https://openreview.net/forum?id=hR_SMu8cxCV.

- Jordan Hoffmann, Sebastian Borgeaud, Arthur Mensch, Elena Buchatskaya, Trevor Cai, Eliza Rutherford, Diego de Las Casas, Lisa Anne Hendricks, Johannes Welbl, Aidan Clark, et al. Training compute-optimal large language models. In *Proceedings of the 36th International Conference on Neural Information Processing Systems*, pp. 30016–30030, 2022.
- Audrey Huang, Adam Block, Dylan J Foster, Dhruv Rohatgi, Cyril Zhang, Max Simchowitz, Jordan T Ash, and Akshay Krishnamurthy. Self-improvement in language models: The sharpening mechanism. *International Conference on Learning Representations (ICLR)*, 2025a.
- Audrey Huang, Adam Block, Qinghua Liu, Nan Jiang, Akshay Krishnamurthy, and Dylan J Foster. Is best-of-n the best of them? coverage, scaling, and optimality in inference-time alignment. *International Conference on Machine Learning (ICML)*, 2025b.
- Audrey Huang, Wenhao Zhan, Tengyang Xie, Jason D Lee, Wen Sun, Akshay Krishnamurthy, and Dylan J Foster. Correcting the mythos of kl-regularization: Direct alignment without overoptimization via chi-squared preference optimization. In *The Thirteenth International Conference on Learning Representations*, 2025c.
- Yuzhen Huang, Jinghan Zhang, Zifei Shan, and Junxian He. Compression represents intelligence linearly. In *First Conference on Language Modeling*, 2024.
- Arthur Jacot, Franck Gabriel, and Clément Hongler. Neural tangent kernel: Convergence and generalization in neural networks. In *Advances in Neural Information Processing Systems (NeurIPS)*, 2018.
- Xiang Ji, Sanjeev Kulkarni, Mengdi Wang, and Tengyang Xie. Self-play with adversarial critic: Provable and scalable offline alignment for language models. *arXiv:2406.04274*, 2024.
- Nan Jiang and Tengyang Xie. Offline reinforcement learning in large state spaces: Algorithms and guarantees. *Statistical Science*, 2024.
- Hangzhan Jin, Sitao Luan, Sicheng Lyu, Guillaume Rabusseau, Reihaneh Rabbany, Doina Precup, and Mohammad Hamdaqa. RL fine-tuning heals ood forgetting in sft, 2025. URL <https://arxiv.org/abs/2509.12235>.
- Ying Jin, Zhuoran Yang, and Zhaoran Wang. Is pessimism provably efficient for offline rl? In *International conference on machine learning*, pp. 5084–5096. PMLR, 2021.
- Jared Kaplan, Sam McCandlish, Tom Henighan, Tom B Brown, Benjamin Chess, Rewon Child, Scott Gray, Alec Radford, Jeffrey Wu, and Dario Amodei. Scaling laws for neural language models. *arXiv preprint arXiv:2001.08361*, 2020.
- Diederik P Kingma and Jimmy Ba. Adam: A method for stochastic optimization. *International Conference on Learning Representations*, 2015.
- Hong Liu, Sang Michael Xie, Zhiyuan Li, and Tengyu Ma. Same Pre-training Loss, Better Downstream: Implicit Bias Matters for Language Models, 2022.
- Zhihan Liu, Miao Lu, Shenao Zhang, Boyi Liu, Hongyi Guo, Yingxiang Yang, Jose Blanchet, and Zhaoran Wang. Provably mitigating overoptimization in RLHF: Your SFT loss is implicitly an adversarial regularizer. *arXiv:2405.16436*, 2024.
- Sanae Lotfi, Marc Finzi, Yilun Kuang, Tim GJ Rudner, Micah Goldblum, and Andrew Gordon Wilson. Non-vacuous generalization bounds for large language models. *arXiv preprint arXiv:2312.17173*, 2023.
- Sanae Lotfi, Yilun Kuang, Marc Finzi, Brandon Amos, Micah Goldblum, and Andrew G Wilson. Unlocking tokens as data points for generalization bounds on larger language models. *Advances in Neural Information Processing Systems*, 37:9229–9256, 2024.
- Nicholas Lourie, Michael Y Hu, and Kyunghyun Cho. Scaling laws are unreliable for downstream tasks: A reality check. *arXiv preprint arXiv:2507.00885*, 2025.

- Xingyu Lu, Xiaonan Li, Qinyuan Cheng, Kai Ding, Xuanjing Huang, and Xipeng Qiu. Scaling laws for fact memorization of large language models. In Yaser Al-Onaizan, Mohit Bansal, and Yun-Nung Chen (eds.), *Findings of the Association for Computational Linguistics: EMNLP 2024*, pp. 11263–11282, Miami, Florida, USA, November 2024. Association for Computational Linguistics. doi: 10.18653/v1/2024.findings-emnlp.658. URL <https://aclanthology.org/2024.findings-emnlp.658/>.
- Shahar Mendelson. Learning without Concentration. In *Conference on Learning Theory*, 2014.
- Shahar Mendelson. Extending the scope of the small-ball method. *arXiv preprint arXiv:1709.00843*, 2017.
- Vaishnavh Nagarajan and J. Zico Kolter. Uniform convergence may be unable to explain generalization in deep learning. In *Advances in Neural Information Processing Systems (NeurIPS)*, 2019.
- Vaishnavh Nagarajan, Chen Henry Wu, Charles Ding, and Aditi Raghunathan. Roll the dice & look before you leap: Going beyond the creative limits of next-token prediction. In *Forty-second International Conference on Machine Learning*, 2025.
- Behnam Neyshabur, Ryota Tomioka, and Nati Srebro. Norm-based capacity control in neural networks. In *Conference on Learning Theory (COLT)*, 2015.
- OpenAI. Introducing openai o1. *Blog*, 2024. URL <https://openai.com/o1/>.
- Alec Radford, Jeffrey Wu, Rewon Child, David Luan, Dario Amodei, Ilya Sutskever, et al. Language models are unsupervised multitask learners. *OpenAI blog*, 1(8):9, 2019.
- Alexander Rakhlin and Karthik Sridharan. Statistical learning and sequential prediction, 2012. Available at http://www.mit.edu/~rakhlin/courses/stat928/stat928_notes.pdf.
- Dhruv Rohatgi, Adam Block, Audrey Huang, Akshay Krishnamurthy, and Dylan J. Foster. Computational-statistical tradeoffs at the next-token prediction barrier: Autoregressive and imitation learning under misspecification. *arXiv preprint arXiv:2502.12465*, 2025.
- Clayton Sanford, Bahare Fatemi, Ethan Hall, Anton Tsitsulin, Mehran Kazemi, Jonathan Halcrow, Bryan Perozzi, and Vahab Mirrokni. Understanding transformer reasoning capabilities via graph algorithms. *Advances in Neural Information Processing Systems*, 37:78320–78370, 2024.
- Abulhair Saparov, Srushti Ajay Pawar, Shreyas Pimpalgaonkar, Nitish Joshi, Richard Yuanzhe Pang, Vishakh Padmakumar, Mehran Kazemi, Najoung Kim, and He He. Transformers struggle to learn to search. In *The Thirteenth International Conference on Learning Representations, ICLR 2025, Singapore, April 24-28, 2025*. OpenReview.net, 2025.
- Nikhil Sardana, Jacob Portes, Sasha Dobov, and Jonathan Frankle. Beyond chinchilla-optimal: Accounting for inference in language model scaling laws. In *Forty-first International Conference on Machine Learning*, 2024. URL <https://openreview.net/forum?id=0bmXrtTDUu>.
- Charlie Victor Snell, Jaehoon Lee, Kelvin Xu, and Aviral Kumar. Scaling LLM test-time compute optimally can be more effective than scaling parameters for reasoning. In *The Thirteenth International Conference on Learning Representations*, 2025. URL <https://openreview.net/forum?id=4FWAwZtd2n>.
- Yuda Song, Gokul Swamy, Aarti Singh, J Bagnell, and Wen Sun. The importance of online data: Understanding preference fine-tuning via coverage. *Advances in Neural Information Processing Systems*, 37:12243–12270, 2024.
- Vladimir Spokoiny. Parametric estimation. finite sample theory. *The Annals of Statistics*, pp. 2877–2909, 2012.
- Jacob Mitchell Springer, Sachin Goyal, Kaiyue Wen, Tanishq Kumar, Xiang Yue, Sadhika Malladi, Graham Neubig, and Aditi Raghunathan. Overtrained language models are harder to fine-tune. In *Forty-second International Conference on Machine Learning*, 2025. URL <https://openreview.net/forum?id=YW6edSufht>.

- Yu Sun, Xiaolong Wang, Zhuang Liu, John Miller, Alexei Efros, and Moritz Hardt. Test-time training with self-supervision for generalization under distribution shifts. In *International conference on machine learning*, pp. 9229–9248. PMLR, 2020.
- Jianheng Tang, Qifan Zhang, Yuhan Li, Nuo Chen, and Jia Li. Grapharena: Evaluating and exploring large language models on graph computation. In *The Thirteenth International Conference on Learning Representations*, 2025.
- Alexander K Taylor, Anthony Cuturrufo, Vishal Yathish, Mingyu Derek Ma, and Wei Wang. Are large-language models graph algorithmic reasoners? *arXiv preprint arXiv:2410.22597*, 2024.
- S. A. van de Geer. *Empirical Processes in M-Estimation*. Cambridge University Press, 2000.
- Aad W Van der Vaart. *Asymptotic statistics*, volume 3. Cambridge University Press, 2000.
- Heng Wang, Shangbin Feng, Tianxing He, Zhaoxuan Tan, Xiaochuang Han, and Yulia Tsvetkov. Can language models solve graph problems in natural language? *Advances in Neural Information Processing Systems*, 36:30840–30861, 2023.
- Xinyi Wang, Shawn Tan, Mingyu Jin, William Yang Wang, Rameswar Panda, and Yikang Shen. Do larger language models imply better generalization? a pretraining scaling law for implicit reasoning, 2025. URL <https://arxiv.org/abs/2504.03635>.
- Wing Hung Wong and Xiaotong Shen. Probability inequalities for likelihood ratios and convergence rates of sieve mles. *The Annals of Statistics*, 1995.
- Fang Wu, Weihao Xuan, Ximing Lu, Zaid Harchaoui, and Yejin Choi. The Invisible Leash: Why RLVR May Not Escape Its Origin, 2025.
- Mengzhou Xia, Mikel Artetxe, Chunting Zhou, Xi Victoria Lin, Ramakanth Pasunuru, Danqi Chen, Luke Zettlemoyer, and Ves Stoyanov. Training trajectories of language models across scales. *arXiv preprint arXiv:2212.09803*, 2022.
- Tengyang Xie and Nan Jiang. Q* approximation schemes for batch reinforcement learning: A theoretical comparison. In *Conference on Uncertainty in Artificial Intelligence*, 2020.
- Yuhong Yang and Andrew R Barron. An asymptotic property of model selection criteria. *IEEE Transactions on Information Theory*, 44(1):95–116, 1998.
- Gilad Yehudai, Noah Amsel, and Joan Bruna. Compositional reasoning with transformers, rnns, and chain of thought. *arXiv preprint arXiv:2503.01544*, 2025.
- Yang Yue, Zhiqi Chen, Rui Lu, Andrew Zhao, Zhaokai Wang, Shiji Song, and Gao Huang. Does reinforcement learning really incentivize reasoning capacity in llms beyond the base model? *arXiv preprint arXiv:2504.13837*, 2025.
- Hansi Zeng, Kai Hui, Honglei Zhuang, Zhen Qin, Zhenrui Yue, Hamed Zamani, and Dana Alon. Can Pre-training Indicators Reliably Predict Fine-tuning Outcomes of LLMs?, 2025.
- Chiyan Zhang, Samy Bengio, Moritz Hardt, Benjamin Recht, and Oriol Vinyals. Understanding deep learning requires rethinking generalization. In *International Conference on Learning Representations (ICLR)*, 2017.
- Tong Zhang. Covering number bounds of certain regularized linear function classes. *Journal of Machine Learning Research*, 2(Mar):527–550, 2002.
- Tong Zhang. From ϵ -entropy to KL-entropy: Analysis of minimum information complexity density estimation. *The Annals of Statistics*, 2006.

CONTENTS OF APPENDIX

I	Additional Discussion and Results	16
A	Discussion and Future Work	16
A.1	Simplifications in the Problem Formulation	16
A.2	Future Work	16
B	Additional Related Work	16
C	Properties of the Coverage Profile	17
D	Supporting Results	19
D.1	Properties of the Coverage Profile	19
D.2	Analysis of Best-of-N Sampling under a Good Coverage Profile	20
D.3	Properties of Maximum Likelihood	21
D.4	Autoregressive Models: Coverage and Stopped KL-Divergence	22
E	Comparison to Classical Generalization Bounds for Maximum Likelihood	23
F	Additional Results	25
F.1	Tightness of Theorem 4.1	25
F.2	Maximum Likelihood: Tighter Rates for Convex Classes	25
F.3	Stochastic Gradient Descent: Improved Gradient Normalization for Distillation	26
F.4	An improved tournament via on-policy generation	26
G	Experiments	28
G.1	Graph Reasoning Task	28
G.2	Experiment Details for Figure 1	29
G.3	Experiment Details for Figure 2	31
II	Proofs	32
H	Technical tools	32
H.1	Concentration inequalities	32
I	Proofs from Section 3	33
J	Proofs from Section 4	34
J.1	Proof Sketch for Theorem 4.1	34
J.2	Proof of Theorem 4.1	34
J.3	Proof of Theorem F.1	37
J.4	Proofs for Supporting Results	38
K	Proofs from Section 5	39
K.1	Proof of Proposition 5.1 (Upper Bound)	39
K.2	Proof of Theorem 5.1	40
K.3	Proof of Theorem 4.2	44
K.4	Proof of Proposition 5.1 (Lower Bound)	49
K.5	Proof of the supporting results	51
L	Proofs from Section 6	52
L.1	Proof of Theorem 6.1	52
L.2	Proof of Theorem F.2	54
L.3	Proofs from Section 6.2 (Selection)	55
L.4	Proof of Theorem F.3	56

Part I

Additional Discussion and Results

A DISCUSSION AND FUTURE WORK

Our work, through the lens of coverage, takes a first step toward clarifying the mechanisms through which pre-training with next-token prediction leads to models for which post-training is effective.

A.1 SIMPLIFICATIONS IN THE PROBLEM FORMULATION

In the course of the paper we have made various simplifying assumptions, some of which can be relaxed in a straightforward fashion, while others are more fundamental:

- In language model pre-training, the pre-training corpus consists of sequences y with varying lengths H , and does not typically split examples into prompts and responses. Our formulation in [Section 2](#) is a simplification (one that is closer in spirit to supervised fine-tuning), but we expect that the insights derived here can extend to the general setting.
- Much of our analysis focuses on the realizable/well-specified setting where $\pi_D \in \Pi$. We give evidence in [Appendix F](#) that the coverage profile is more tolerant to misspecification than KL-divergence, but we leave a deeper investigation for future work.
- Our treatment assumes the distribution over prompts μ is the same for pre-training and post-training. This is straightforward to relax at the cost of introducing an additional coverage or distribution shift coefficient to handle the mismatch between the two distributions.
- We show that a good coverage profile is necessary for BoN to succeed on downstream tasks. While there is ample evidence current RL techniques can fail in the absence of coverage ([Yue et al., 2025](#); [Gandhi et al., 2025](#); [Wu et al., 2025](#)), it is not clear what the minimal conditions required for RL are, and they may be weaker than coverage.

A.2 FUTURE WORK

Our results open several new directions for future research.

- *Interventions for coverage.* There is much to be done in understanding and improving existing algorithms such as optimizers through the lens of coverage. Our results in [Section 6](#) show initial promise for using coverage to guide design of optimizers and model selection schemes, but the algorithm design space remains opaque, and there may be significant room for further improvement. More ambitiously, one could imagine re-structuring the entire language modeling pipeline itself around coverage.
- *Semantic coverage.* The notion of coverage we focus on, the *coverage profile*, is mathematically convenient but may be conservative in regard to downstream performance, since it only depends on the model through its predicted probabilities. An important direction for future work is to understand pre-training and post-training through fine-grained “semantic” notions of coverage that more explicitly account for the representations learned by next-token prediction.

B ADDITIONAL RELATED WORK

Related empirical observations. On the empirical side, our results are connected to a line of work that studies to scaling laws for zero-shot downstream performance based on pre-training metrics such as cross-entropy ([Gadre et al., 2024](#); [Huang et al., 2024](#); [Chen et al., 2024](#); [Sardana et al., 2024](#)). Several empirical works have also investigated how specific capabilities scale with additional pre-training, including machine translation ([Ghorbani et al., 2022](#)), knowledge capacity and memorization ([Allen-Zhu & Li, 2025](#); [Lu et al., 2024](#)), and multi-hop reasoning ([Wang et al., 2025](#)). Our findings are consistent with [Liu et al. \(2022\)](#); [Zeng et al. \(2025\)](#); [Lourie et al. \(2025\)](#); [Springer et al. \(2025\)](#), who observe that cross-entropy is not always sufficient for predicting downstream performance, and in some cases can be anti-correlated.

Perhaps most closely related, [Chen et al. \(2025\)](#) show empirically the decreasing cross-entropy in pre-training does not necessarily lead to better pass@N performance, and that pass@N can even degrade as pre-training proceeds—a finding similar to [Figure 1](#).⁵ Our theoretical results can be viewed as placing their findings on stronger theoretical footing; conversely, their empirical results provide strong motivation for our theoretical treatment. [Chen et al. \(2025\)](#) also study a modification to the maximum likelihood objective aimed at improving coverage (in the spirit of [Section 6](#)), but, when instantiated with chain-of-thought, their approach requires a small space of possible final answers.

We mention in passing a few additional works. [Chu et al. \(2025\)](#) explored the different synergistic roles that supervised fine-tuning (SFT) and RL play in language model development, and subsequent work observed that the best checkpoint to start RL from can sometimes be in the middle of SFT ([Jin et al., 2025](#)). [Bansal et al. \(2025\)](#) empirically identified the coverage of teacher-generated synthetic data as an important indicator of how effective distillation would be for reasoning tasks. Several papers have also investigated empirical tradeoffs between model size and reasoning performance under best-of-N sampling ([Snell et al., 2025](#); [Brown et al., 2025](#)).

Coverage in post-training. Coverage metrics similar to coverage profile play a central role in theoretical literature on post-training and test-time algorithms ([Huang et al., 2025a;b;c](#); [Foster et al., 2025](#); [Liu et al., 2024](#); [Song et al., 2024](#); [Gao et al., 2024](#); [Liu et al., 2024](#); [Ji et al., 2024](#)), which analyze algorithms under the assumption that the base model has good coverage; our work can be viewed as providing theoretical motivation for this assumption.

Various notions of coverage similar to coverage profile have also appeared in the more classical literature on offline reinforcement learning ([Farahmand et al., 2010](#); [Chen & Jiang, 2019](#); [Xie & Jiang, 2020](#); [Jin et al., 2021](#); [Foster et al., 2022](#); [Jiang & Xie, 2024](#)); here coverage is typically used to quantify the quality of an offline dataset rather than a model/policy itself.

Generalization in deep learning. Understanding the generalization behavior of deep learning models has been a central focus of the theory community for the last decade ([Neyshabur et al., 2015](#); [Zhang et al., 2017](#); [Bartlett et al., 2017](#); [Jacot et al., 2018](#); [Belkin et al., 2019](#); [Nagarajan & Kolter, 2019](#); [Bartlett et al., 2020](#); [Bartlett & Montanari, 2021](#)). Our approach is somewhat complementary, in the sense that it focuses on the specific objective of next-token prediction with the logarithmic loss, and aims to understand when minimizing this loss leads to generalization for an *alternative* objective, coverage profile. We expect that our techniques can be combined with these more general results to provide more refined understanding of generalization for coverage profile with deep models.

From this line of work, perhaps most closely related are [Lotfi et al. \(2023; 2024\)](#); [Finzi et al. \(2025\)](#), which aim to provide non-vacuous generalization bounds for the cross-entropy loss itself for autoregressive models.

Analysis of maximum likelihood. On the theoretical side, our results are most closely related to a classical line of work in statistics ([Wong & Shen, 1995](#); [van de Geer, 2000](#); [Zhang, 2006](#)), which shows that maximum likelihood can converge to the true model in Hellinger distance (or other Renyi divergences) under minimal assumptions, even when KL divergence is poorly behaved (large or infinite). Our results in [Section 4](#) are similar in spirit, but provide a more fine-grained perspective, showing that coverage profile can converge even faster than these results might suggest, particularly as one ventures further into the tail. Our analysis has some conceptual similarity to the small ball method of [Mendelson \(2014; 2017\)](#), which we elaborate on in [Appendix J.1](#).

Our techniques are also related to recent work of [Foster et al. \(2024\)](#); [Rohatgi et al. \(2025\)](#), which specialize the general techniques above to autoregressive models (e.g., under Hellinger distance).

C PROPERTIES OF THE COVERAGE PROFILE

Before proceeding, we briefly discuss some conceptual properties of the coverage profile that will be helpful to keep in mind.

Remark C.1 (Coverage profile as a refinement of cross-entropy). *While we position the coverage profile as a new quantity of interest, it can also be viewed as a fine-grained, inference budget-sensitive*

⁵While [Chen et al. \(2025\)](#) use the term “coverage”, it is used as a synonym for pass@N, and is not specifically related to the notion of the coverage profile we consider here.

refinement of cross-entropy. Concretely, if we write

$$\text{Cov}_N(\pi_D \parallel \hat{\pi}) = \mathbb{P}_{\pi_D} \left[\log \frac{\pi_D(y \mid x)}{\hat{\pi}(y \mid x)} \geq \log N \right], \quad (15)$$

it becomes clear that the coverage profile is simply the cumulative distribution function (CDF) of the log density ratio $X := \log \frac{\pi_D(y \mid x)}{\hat{\pi}(y \mid x)}$, while KL-divergence corresponds to the mean: $\mathbb{E}_{\pi_D}[X]$. It is well known that the CDF of a random variable is a more informative statistic than its mean (Durrett, 2019); the former can be much more sensitive to the model’s behavior at the tail than the latter. Indeed, the coverage profile can behave very differently across scales, as shown by Figure 1.⁶

Remark C.2 (KL divergence and coverage profile are not estimable). We emphasize that KL-divergence and the coverage profile are not estimable quantities in general, due to the fact both depend on the unknown density $\pi_D(y \mid x)$ for the data distribution. This motivates the use of cross-entropy in practice, as the former is an estimable upper bound on $D_{\text{KL}}(\pi_D \parallel \hat{\pi})$. Analogously, we show in Section 6.2 that various estimable proxies for the coverage profile can be used to select models with good coverage.

An exception is the expert distillation setting (see Section 6.1), where π_D is a teacher network for which the log-probabilities $\log \pi_D(y \mid x)$ are available.

Remark C.3 (Sequence-level versus answer-level coverage). Our discussion so far has focused on coverage at the sequence level. For reasoning tasks, it is natural to explicitly factorize the response $y = (y_{\text{cot}}, y_{\text{ans}})$ into a chain-of-thought (reasoning trajectory) component y_{cot} and an answer component y_{ans} . For this setting, a weaker notion coverage is the following answer-level coverage profile:

$$\text{Cov}_N^{\text{ans}}(\pi_D \parallel \hat{\pi}) := \mathbb{P}_{\pi_D} \left[\frac{\pi_D(y_{\text{ans}} \mid x)}{\hat{\pi}(y_{\text{ans}} \mid x)} \geq N \right].$$

Informally, the answer-level coverage profile is sufficient for downstream BoN success for tasks where it is only important to produce the right answer, not a correct reasoning trace. We have $\text{Cov}_N^{\text{ans}}(\pi_D \parallel \hat{\pi}) \leq \text{Cov}_N(\pi_D \parallel \hat{\pi})$, but the former can be strictly smaller in general.

We hope that by providing a comprehensive understanding of sequence-level coverage, our work can set the stage for future research on answer-level coverage and other finer-grained notions of coverage; we give some initial results along these lines in Appendix F.

⁶Interestingly, we show (Proposition D.1) that if the coverage profile satisfies a certain growth condition uniformly for all scales M , then it implies a bound on KL-divergence—a weak converse to Proposition 3.1.

D SUPPORTING RESULTS

D.1 PROPERTIES OF THE COVERAGE PROFILE

Proposition D.1 (KL-to-coverage conversion). *For all models π_D and π and $M \geq 2$, we have*

$$\text{Cov}_N(\pi) \leq \frac{D_{\text{KL}}(\pi_D \| \pi)}{\log N - 1 + \frac{1}{N}}.$$

Proof of Proposition D.1. Lemma 27 of [Block & Polyanskiy \(2023\)](#) states that for any $N > 1$ and any convex $f : [0, \infty] \rightarrow [0, \infty]$ with $f(1) = f'(1) = 0$,

$$\text{Cov}_N(\pi) = \mathbb{P}_{\pi_D} \left[\frac{\pi_D(y|x)}{\pi(y|x)} > N \right] \leq \frac{ND_f(\pi_D \| \pi)}{f(N)}, \quad (16)$$

where $D_f(\pi_D \| \pi) := \mathbb{E}_{\pi} \left[f \left(\frac{d\pi_D}{d\pi} \right) \right]$. Applying this with KL-divergence, which corresponds to $f(x) = x \log x - x + 1$ with $f'(x) = \log x$, we have that

$$\frac{N}{f(N)} = \frac{1}{\log N - 1 + 1/N}, \quad (17)$$

which gives the result. $\square$

Proposition D.2 (Tightness of KL-to-coverage conversion). *For any $N \geq 2$, there exist models π_D and $\hat{\pi}$ such that*

$$\text{Cov}_N(\hat{\pi}) \geq \frac{D_{\text{KL}}(\pi_D \| \hat{\pi})}{\log N - \frac{1}{2} + \frac{1}{2N}}.$$

Proof of Proposition D.2. Consider $\pi_D = \text{Ber}(p)$ and $\hat{\pi} = \text{Ber}(p/N)$ with $p \leq \frac{1}{2}$. Then $\text{Cov}_N(\hat{\pi}) = p$ and

$$\begin{aligned} D_{\text{KL}}(\pi_D \| \hat{\pi}) &= p \log N + (1-p) \log \frac{1-p}{1-\frac{p}{N}} \leq p \log N + (1-p) \left(\frac{1-p}{1-\frac{p}{N}} - 1 \right) \\ &= p \left(\log N - (1-p) \frac{1-\frac{1}{N}}{1-\frac{p}{N}} \right) \\ &\leq p \cdot \left(\log N - \frac{1}{2} + \frac{1}{2N} \right). \end{aligned}$$

This is the desired result. $\square$

Proposition D.3 (Uniform coverage decay implies bounded KL). *Given $\pi, \pi_D : \mathcal{X} \rightarrow \Delta(\mathcal{Y})$, define $W_{\max} := \sup_{x,y} \frac{\pi_D(y|x)}{\pi(y|x)}$ and*

$$C := \sup_{N \geq 1} \{ \text{Cov}_N(\pi) \cdot \log N \},$$

where we note that $C \leq \log W_{\max}$. It holds that

$$D_{\text{KL}}(\pi_D \| \pi) \leq C \cdot (1 + \log(\log(W_{\max})/C)). \quad (18)$$

Proof of Proposition D.3. Let $\delta > 0$ a fixed parameter, and define $X := \pi_D/\pi$. Then we have

$$D_{\text{KL}}(\pi_D \| \pi) = \mathbb{E}_{\pi_D}[\log(X)] \leq \mathbb{E}_{\pi_D}[\log(X) \mathbb{I}\{\log(X) > \delta\}] + \delta. \quad (19)$$

Since $X \leq W_{\max}$ almost surely, we can write

$$\mathbb{E}_{\pi_D}[\log(X)\mathbb{I}\{\log(X) > \delta\}] = \int_{\delta}^{\log(W_{\max})} \mathbb{P}_{\pi_D}[\log(X) > t] dt \quad (20)$$

$$= \int_{\delta}^{\log(W_{\max})} \mathbb{P}_{\pi_D}[X > e^t] dt \quad (21)$$

$$\leq C \int_{\delta}^{\log(W_{\max})} \frac{1}{t} dt \quad (22)$$

$$= C \log\left(\frac{\log(W_{\max})}{\delta}\right). \quad (23)$$

The result now follows by setting $\delta = C$. $\square$

Proposition D.4 (Hellinger-to-coverage conversion). *For all models π_D and π and $N > 1$, we have*

$$\text{Cov}_N(\pi_D \parallel \pi) \leq \frac{2N}{(\sqrt{N} - 1)^2} \cdot D_H^2(\pi_D, \pi).$$

Proof of Proposition D.4. Without loss of generality, we assume $\mathcal{Y}$ is discrete in the following proof. By definition,

$$\begin{aligned} D_H^2(\pi_D, \pi) &= \frac{1}{2} \mathbb{E}_{x \sim \pi_D} \left[\sum_y \left(\sqrt{\pi_D(y \mid x)} - \sqrt{\pi(y \mid x)} \right)^2 \right] \\ &\geq \frac{1}{2} \mathbb{E}_{x \sim \pi_D} \left[\sum_y \pi_D(y \mid x) \left(1 - \frac{1}{\sqrt{N}} \right)^2 \mathbb{I} \left\{ \pi(y \mid x) \leq \frac{1}{N} \pi_D(y \mid x) \right\} \right] \\ &= \frac{1}{2} \left(1 - \frac{1}{\sqrt{N}} \right)^2 \mathbb{P}_{\pi_D} \left[\frac{\pi_D(y \mid x)}{\pi(y \mid x)} > N \right], \end{aligned}$$

where the inequality follows from the fact that $\sqrt{\pi_D(y \mid x)} - \sqrt{\pi(y \mid x)} \geq \left(1 - \frac{1}{\sqrt{N}} \right) \sqrt{\pi_D(y \mid x)}$ is implied by $\pi(y \mid x) \leq \frac{1}{N} \pi_D(y \mid x)$. Re-organizing completes the proof. $\square$

Proposition D.5 (Chain rule for coverage profile). *For any models π_D, π_T , and $\hat{\pi}$, and any $M_1, M_2 \geq 2$, we have*

$$\text{Cov}_{M_1}(\pi_T \parallel \hat{\pi}) \leq M_2 \cdot \text{Cov}_{M_1/M_2}(\pi_D \parallel \hat{\pi}) + \text{Cov}_{M_2}(\pi_T \parallel \pi_D). \quad (24)$$

Proof of Proposition D.5. We can write

$$\begin{aligned} \text{Cov}_{M_1}(\pi_T \parallel \hat{\pi}) &= \mathbb{P}_{\pi_T} \left[\frac{\pi_T(y \mid x)}{\hat{\pi}(y \mid x)} > M_1 \right] \\ &= \mathbb{P}_{\pi_T} \left[\frac{\pi_T(y \mid x)}{\hat{\pi}(y \mid x)} > M_1, \frac{\pi_T(y \mid x)}{\pi_D(y \mid x)} \leq M_2 \right] + \mathbb{P}_{\pi_T} \left[\frac{\pi_T(y \mid x)}{\hat{\pi}(y \mid x)} > M_1, \frac{\pi_T(y \mid x)}{\pi_D(y \mid x)} > M_2 \right] \\ &\leq M_2 \mathbb{P}_{\pi_D} \left[\frac{\pi_D(y \mid x)}{\hat{\pi}(y \mid x)} > M_1/M_2 \right] + \mathbb{P}_{\pi_T} \left[\frac{\pi_T(y \mid x)}{\pi_D(y \mid x)} > M_2 \right] \\ &= M_2 \text{Cov}_{M_1/M_2}(\pi_D \parallel \hat{\pi}) + \text{Cov}_{M_2}(\pi_T \parallel \pi_D). \end{aligned}$$

$\square$

D.2 ANALYSIS OF BEST-OF-N SAMPLING UNDER A GOOD COVERAGE PROFILE

In this section we analyze the performance of the Best-of-N algorithm under a good coverage profile. Let a base model $\hat{\pi}$ be given, and let a reward function $r_T(x, y) \in [0, 1]$ be given. Let $\pi_T : \mathcal{X} \rightarrow \Delta(\mathcal{Y})$ denote an arbitrary task-specific comparator policy.

We let $\hat{\pi}_N^{\text{BoN}}(x)$ denote the distribution of the Best-of-N algorithm with parameter N , which draws N responses $y^1, \dots, y^N \stackrel{\text{i.i.d.}}{\sim} \hat{\pi}(\cdot \mid x)$ and returns $y = \arg \max_{y_i} r_T(x, y_i)$.

Proposition D.6 (Coverage implies success for BoN). *Let $M \geq 1$ be given. For any $\varepsilon > 0$, if $N \geq 2M \log(\varepsilon^{-1})$ and $\text{Cov}_M(\pi_{\text{T}} \parallel \hat{\pi}) \leq \frac{1}{2}$, then we are guaranteed that*

$$\mathbb{E}_{x \sim \mu} [r_{\text{T}}(x, \pi_{\text{T}}(x)) - r_{\text{T}}(x, \hat{\pi}_N^{\text{BoN}}(x))] \leq \text{Cov}_M(\pi_{\text{T}} \parallel \hat{\pi}) + \varepsilon. \quad (25)$$

Proof of Proposition D.6. This is an immediate consequence of Lemma F.1 in Huang et al. (2025b), noting that we can bound $\mathcal{E}_M(\pi_{\text{T}} \parallel \hat{\pi}) \leq \text{Cov}_M(\pi_{\text{T}} \parallel \hat{\pi})$. $\square$

Proposition D.7 (Coverage is necessary for BoN). *For any model $\hat{\pi}$ and reference π_{T} , and for any $N \geq 2$, there exists a reward function $r_{\text{T}}(x, y) \in \{0, 1\}$ such that*

$$\mathbb{E}_{x \sim \mu} [r_{\text{T}}(x, \pi_{\text{T}}(x)) - r_{\text{T}}(x, \hat{\pi}_N^{\text{BoN}}(x))] \geq \frac{1}{2} \text{Cov}_{2N}(\pi_{\text{T}} \parallel \hat{\pi}). \quad (26)$$

Proof of Proposition D.7. For any $x \in \mathcal{X}$, we define $S_x := \{y \in \mathcal{Y} : \frac{\pi_{\text{T}}(y|x)}{\hat{\pi}(y|x)} \geq 2N\}$ and let $r_{\text{T}}(x, y) = \mathbb{I}\{y \in S_x\}$.

By definition, for any fixed $x \in \mathcal{X}$, it holds that

$$\begin{aligned} r_{\text{T}}(x, \hat{\pi}_N^{\text{BoN}}(x)) &= \mathbb{P}_{y \sim \hat{\pi}_N^{\text{BoN}}(x)}(y \in S_x) = \mathbb{P}_{y^1, \dots, y^N \sim \hat{\pi}(\cdot|x)}(\exists i \in [N], y^i \in S_x) \\ &= 1 - (1 - \mathbb{P}_{y \sim \hat{\pi}(\cdot|x)}(y \in S_x))^N \leq N \cdot \mathbb{P}_{y \sim \hat{\pi}(\cdot|x)}(y \in S_x) \\ &= N \cdot \sum_{y \in S_x} \hat{\pi}(y | x) \leq N \cdot \sum_{y \in S_x} \frac{1}{2N} \pi_{\text{T}}(y | x) = \frac{1}{2} \mathbb{P}_{y \sim \pi_{\text{T}}(\cdot|x)}(S_x), \end{aligned}$$

where we use the fact that $\hat{\pi}(y | x) \leq \frac{1}{2N} \pi_{\text{T}}(y | x)$ for any $y \in S_x$. We also note that $\mathbb{P}_{x \sim \mu, y \sim \pi_{\text{T}}(\cdot|x)}(y \in S_x) = \text{Cov}_{2N}(\pi_{\text{T}} \parallel \hat{\pi})$. Therefore,

$$\mathbb{E}_{x \sim \mu} [r_{\text{T}}(x, \pi_{\text{T}}(x)) - r_{\text{T}}(x, \hat{\pi}_N^{\text{BoN}}(x))] \geq \frac{1}{2} \text{Cov}_{2N}(\pi_{\text{T}} \parallel \hat{\pi}). \quad \square$$

D.3 PROPERTIES OF MAXIMUM LIKELIHOOD

Proposition D.8 (Convergence of maximum likelihood in Hellinger distance). *Assume that $\pi_{\text{D}} \in \Pi$. With probability at least $1 - \delta$, the maximum likelihood estimator $\hat{\pi} := \arg \max_{\pi \in \Pi} \hat{L}_n(\pi)$ satisfies,*

$$D_{\text{H}}^2(\pi_{\text{D}}, \hat{\pi}) \lesssim \inf_{\varepsilon > 0} \left\{ \frac{\log \mathcal{N}_{\infty}(\Pi, \varepsilon)}{n} + \varepsilon \right\}, \quad (27)$$

and consequently

$$\text{Cov}_M(\hat{\pi}) \lesssim \inf_{\varepsilon > 0} \left\{ \frac{\log \mathcal{N}_{\infty}(\Pi, \varepsilon)}{n} + \varepsilon \right\}. \quad (28)$$

for all $M \geq 2$.

Proof of Proposition D.8. The first bound follows from Proposition B.2 of Foster et al. (2024). The second bound follows from applying Proposition D.4. $\square$

Proposition D.9 (Convergence of maximum likelihood in KL). *Assume that $\pi_{\text{D}} \in \Pi$, and that all $\pi \in \Pi$ satisfy $\left\| \frac{\pi_{\text{D}}}{\pi} \right\|_{\infty} \leq W_{\max}$. With probability at least $1 - \delta$, the maximum likelihood estimator $\hat{\pi} := \arg \max_{\pi \in \Pi} \hat{L}_n(\pi)$ satisfies,*

$$D_{\text{KL}}(\pi_{\text{D}} \parallel \hat{\pi}) \lesssim \log W_{\max} \cdot \inf_{\varepsilon > 0} \left\{ \frac{\log \mathcal{N}_{\infty}(\Pi, \varepsilon)}{n} + \varepsilon \right\}, \quad (29)$$

and consequently

$$\text{Cov}_M(\hat{\pi}) \lesssim \frac{\log W_{\max}}{\log M} \cdot \inf_{\varepsilon > 0} \left\{ \frac{\log \mathcal{N}_{\infty}(\Pi, \varepsilon)}{n} + \varepsilon \right\}, \quad (30)$$

for all $M \geq 2$.

We remark that the $\log(W_{\max})$ -factor in Eq. (29) can be tight in general. For example, for the class Π considered in Proposition 3.2, it holds that $\log \mathcal{N}_{\infty}(\Pi, \varepsilon) \lesssim \log(1/\varepsilon) \vee 1$ and $\|\frac{\pi_D}{\pi}\|_{\infty} \leq e^{2H}$.

Proof of Proposition D.9. By Lemma 4 of Yang & Barron (1998), it holds that

$$D_{\text{KL}}(\pi_D \parallel \hat{\pi}) \leq (2 + \log(W_{\max})) D_H^2(\pi_D, \hat{\pi}).$$

Therefore, the first bound then follows from Eq. (27). The second bound follows from applying Proposition D.1. $\square$

D.4 AUTOREGRESSIVE MODELS: COVERAGE AND STOPPED KL-DIVERGENCE

Proposition D.10. *Define*

$$D_{\text{seq}, N}(\pi_D \parallel \pi) = \mathbb{E}_{(x, y_{1:H}) \sim \pi_D} \min \left\{ \log N, \sum_{h=1}^H D_{\text{KL}}(\pi_D(\cdot \mid x, y_{1:h-1}) \parallel \pi(\cdot \mid x, y_{1:h-1})) \right\}. \quad (31)$$

Then for $N > e$, it holds that

$$\text{Cov}_N(\pi_D \parallel \pi) \leq \frac{2}{\log N - 1} D_{\text{seq}, N}(\pi_D \parallel \pi). \quad (32)$$

Proof of Proposition D.10. Consider the stopping time

$$\tau := \min \left\{ h : h = H \text{ or } \sum_{j \leq h} D_{\text{KL}}(\pi_D(y_{j+1} = \cdot \mid x, y_{1:j}) \parallel \pi(y_{j+1} = \cdot \mid x, y_{1:j})) > \log N \right\}.$$

Then, for the process $Y^\tau = (x, y_{1:\tau})$, we have the chain rule:

$$\begin{aligned} & D_{\text{KL}}(\pi_D(Y^\tau = \cdot) \parallel \pi(Y^\tau = \cdot)) \\ &= \mathbb{E}_{\pi_D} \left[\sum_{h=1}^{\tau} D_{\text{KL}}(\pi_D(y_h = \cdot \mid x, y_{1:h-1}) \parallel \pi(y_h = \cdot \mid x, y_{1:h-1})) \right] \\ &\leq \mathbb{E}_{\pi_D} \min \left\{ \log N, \sum_{h=1}^H D_{\text{KL}}(\pi_D(y_h = \cdot \mid x, y_{1:h-1}) \parallel \pi(y_h = \cdot \mid x, y_{1:h-1})) \right\}, \end{aligned}$$

where the inequality uses $\sum_{j < \tau} D_{\text{KL}}(\pi_D(y_{j+1} = \cdot \mid x, y_{1:j}) \parallel \pi(y_{j+1} = \cdot \mid x, y_{1:j})) \leq \log N$, which follows from the definition of τ . Therefore, by Proposition D.1, we have

$$\mathbb{P}_{\pi_D} \left(\frac{\pi_D(Y^\tau)}{\pi(Y^\tau)} \geq \log N \right) \leq \frac{D_{\text{KL}}(\pi_D(Y^\tau = \cdot) \parallel \pi(Y^\tau = \cdot))}{\log N - 1 + 1/N}.$$

Finally, we bound

$$\mathbb{P}_{\pi_D} \left(\frac{\pi_D(y_{1:H} \mid x)}{\pi(y_{1:H} \mid x)} \geq N \right) \leq \mathbb{P}_{\pi_D}(\tau < H) + \mathbb{P}_{\pi_D} \left(\frac{\pi_D(Y^\tau)}{\pi(Y^\tau)} \geq \log N \right).$$

By Markov's inequality,

$$\begin{aligned} \mathbb{P}_{\pi_D}(\tau < H) &\leq \mathbb{P}_{\pi_D} \left(\sum_{h=1}^H D_{\text{KL}}(\pi_D(\cdot \mid x, y_{1:h-1}) \parallel \pi(\cdot \mid x, y_{1:h-1})) > \log N \right) \\ &\leq \frac{1}{\log N} \mathbb{E}_{\pi_D} \min \left\{ \log N, \sum_{h=1}^H D_{\text{KL}}(\pi_D(\cdot \mid x, y_{1:h-1}) \parallel \pi(\cdot \mid x, y_{1:h-1})) \right\}. \end{aligned}$$

Combining the inequalities above completes the proof. $\square$

Proposition D.11. *For any $N \geq 1, \delta \in (0, 1)$, it holds that*

$$\text{Cov}_N(\pi_D \parallel \pi) \geq \mathbb{P}_{\pi_D} \left(\sum_{h=1}^H D_H^2(\pi_D(\cdot \mid x, y_{1:h-1}), \pi(\cdot \mid x, y_{1:h-1})) \geq \log(N/\delta) \right) - \delta.$$

Proof of Proposition D.11. By definition,

$$\begin{aligned} & \mathbb{E}_{y_h \sim \pi_D(\cdot | x, y_{1:h-1})} \exp\left(-\frac{1}{2} \log \frac{\pi_D(y_h | x, y_{1:h-1})}{\pi(y | x, y_{1:h-1})}\right) \\ &= \sum_{y_h \in \mathcal{Y}} \sqrt{\pi_D(y_h | x, y_{1:h-1}) \cdot \pi(y | x, y_{1:h-1})} \\ &= 1 - D_H^2(\pi_D(\cdot | x, y_{1:h-1}), \pi(\cdot | x, y_{1:h-1})) \leq \exp(-D_H^2(\pi_D(\cdot | x, y_{1:h-1}), \pi(\cdot | x, y_{1:h-1}))). \end{aligned}$$

Therefore, it holds that

$$\mathbb{E}_{\pi_D} \exp\left(\sum_{h=1}^H D_H^2(\pi_D(\cdot | x, y_{1:h-1}), \pi(\cdot | x, y_{1:h-1})) - \frac{1}{2} \log \frac{\pi_D(y_h | x, y_{1:h-1})}{\pi(y | x, y_{1:h-1})}\right) \leq 1.$$

By Markov inequality, this implies

$$\mathbb{P}_{\pi_D} \left(\frac{1}{2} \log \frac{\pi_D(y_{1:H} | x)}{\pi(y_{1:H} | x)} \leq \sum_{h=1}^H D_H^2(\pi_D(\cdot | x, y_{1:h-1}), \pi(\cdot | x, y_{1:h-1})) - \log(1/\delta) \right) \leq \delta.$$

To conclude, we note that

$$\begin{aligned} & \mathbb{P}_{\pi_D} \left(\sum_{h=1}^H D_H^2(\pi_D(\cdot | x, y_{1:h-1}), \pi(\cdot | x, y_{1:h-1})) \geq \log(N/\delta) \right) \\ & \leq \mathbb{P}_{\pi_D} \left(\sum_{h=1}^H D_H^2(\pi_D(\cdot | x, y_{1:h-1}), \pi(\cdot | x, y_{1:h-1})) \geq \frac{1}{2} \log \frac{\pi_D(y_{1:H} | x)}{\pi(y_{1:H} | x)} + \log(1/\delta) \right) \\ & \quad + \mathbb{P}_{\pi_D} \left(\frac{1}{2} \log \frac{\pi_D(y_{1:H} | x)}{\pi(y_{1:H} | x)} + \log(1/\delta) \geq \log(N/\delta) \right) \\ & \leq \delta + \text{Cov}_N(\pi_D \| \pi). \end{aligned}$$

Re-organizing gives the desired result. $\square$

E COMPARISON TO CLASSICAL GENERALIZATION BOUNDS FOR MAXIMUM LIKELIHOOD

In this section we briefly compare our main coverage-based generalization bound for maximum likelihood to classical generalization bounds for maximum likelihood based on Hellinger distance and KL-divergence.

Comparison to KL concentration. For general model classes Π , the best KL generalization bound we are aware of is [Proposition D.9 \(Appendix D\)](#), which scales as roughly

$$D_{\text{KL}}(\pi_D \| \hat{\pi}) \lesssim \log W_{\max} \cdot \mathcal{C}_{\text{fine}}(\Pi, n)$$

under the assumption that all $\pi \in \Pi$ obey a sequence-level density ratio bound $\|\frac{\pi_D}{\pi}\|_{\infty} \leq W_{\max}$, where $\log \mathcal{N}$ is an appropriate notion of covering number; note that for the autoregressive linear class, we have $\log W_{\max} = BH$, matching [Proposition 3.2](#). Combining such a guarantee with [Proposition 3.1](#) gives a coverage bound of roughly

$$\text{Cov}_N(\hat{\pi}) \lesssim \frac{\log W_{\max}}{\log N} \cdot \mathcal{C}_{\text{fine}}(\Pi, n);$$

this is rather uninteresting since $\text{Cov}_N(\hat{\pi}) = 0$ for $N \geq W_{\max}$; in other words, we do not get a meaningful improvement as we scale N .

Asymptotic bounds for maximum likelihood. We also note that the classical theory of maximum likelihood (e.g., [Van der Vaart \(2000\)](#)) provides the following *asymptotic* convergence rate for d -dimensional parametric classes Π :

$$D_{\text{KL}}(\pi_D \| \hat{\pi}) \lesssim \frac{d}{n} \lesssim \mathcal{C}_{\text{fine}}(\Pi, n), \quad n \rightarrow +\infty.$$

While this upper bound does *not* scale with $\log W_{\max}$, it can only be attained with $n \geq n_0$ for a sufficiently large burn-in cost n_0 (typically scaling with $\log W_{\max}$ itself or similar problem-dependent parameters; see, e.g., [Spokoiny \(2012\)](#) for non-asymptotic bounds of this type).

Comparison to Hellinger concentration. For general model classes Π , the best Hellinger generalization bound we are aware of is [Proposition D.8 \(Appendix D\)](#), which scales as roughly

$$D_{\text{H}}^2(\pi_{\mathbb{D}}, \hat{\pi}) \lesssim \mathcal{C}_{\text{fine}}(\Pi, n)$$

Combining such a guarantee with [Proposition 3.1](#) gives a coverage bound of roughly

$$\text{Cov}_N(\hat{\pi}) \lesssim \mathcal{C}_{\text{fine}}(\Pi, n)$$

for all $N \geq 2$. Compare to the KL-based result above, this result gives a non-trivial bound on coverage when N is constant (comparable to [Theorem 4.1](#)), but the issue is that it gives no further improvement as we scale N .

F ADDITIONAL RESULTS

F.1 TIGHTNESS OF [THEOREM 4.1](#)

To conclude, we show that the coarse and fine-grained terms in [Theorem 4.1](#) are both tight in general.

Proposition F.1. *The following lower bounds on coverage hold for the maximum likelihood estimator.*

(a) Coarse rate: *For any $n, d \geq 1$ and $B \geq \log(5n)$, there exists a class Π with $\log \mathcal{N}_\infty(\Pi, \alpha) \lesssim d \log(B/\alpha) \vee 1$ and $\pi_0 \in \Pi$ such that with probability at least 0.5, it holds that for any $N \leq e^B$,*

$$\text{Cov}_N(\hat{\pi}) \geq c \cdot \frac{d}{n}.$$

(b) Fine rate: *For any $n \geq d \geq 1, N \geq 1$, there exists a class Π and $\pi_0 \in \Pi$ such that $|\Pi| = 2^d + 1$ and $\mathcal{N}_\infty(\Pi, \alpha) \leq 2$ for any $\alpha \geq \sqrt{\frac{d}{n}}$, and with probability at least 0.5, it holds that*

$$\text{Cov}_N(\hat{\pi}) \geq c \cdot \frac{d}{n \cdot \log N}.$$

Informally, case (a) shows that for the class Π under consideration, the coverage does not decrease with $\log N$ until N is trivially large such that $\log \mathcal{N}_\infty(\Pi, \log N) = 0$; this is precisely the behavior of the coarse term in [Theorem 4.1](#), so this implies there is no hope of removing this term. Meanwhile, case (b) can be interpreted as showing that there is no hope of replacing the high-precision covering number found in the fine-grained term in [Theorem 4.1](#) with a coarser notion (e.g. at the scale in the coarse-grained term), since the rate grows with $d \approx \log |\Pi|$ even though $\log \mathcal{N}(\Pi, \alpha)$ is constant for $\alpha \geq \sqrt{\frac{d}{n}}$. We note that [Proposition F.1](#) is an algorithm-specific lower bound, not an information-theoretic lower bound; we show in [Section 6.2](#) is that it is possible to improve over [Theorem 4.1](#) with algorithms explicitly designed to optimize for coverage.

F.2 MAXIMUM LIKELIHOOD: TIGHTER RATES FOR CONVEX CLASSES

In the following, we analyze the MLE for *convex* model class.

Assumption F.1 (Convex model class). *The class Π satisfies $\Pi = \{\pi_\theta : \theta \in \Theta\}$ for a convex, compact parameter space Θ , and the mapping $\theta \mapsto \pi_\theta(y | x)$ is concave for all $x \in \mathcal{X}, y \in \mathcal{Y}$.*

Theorem F.1 (Fast convergence of coverage for convex classes). *Let $\alpha \geq 0, N' \geq 1, N \geq 2e^{2\alpha} N'$ be given, and suppose that [Assumption F.1](#) holds. Define*

$$\theta^* = \arg \min_{\theta \in \Theta} D_{\text{KL}}(\pi_0 \| \pi_\theta).$$

With probability at least $1 - \delta$, the maximum likelihood estimator $\hat{\pi} := \arg \max_{\pi \in \Pi} \hat{L}_n(\pi)$ satisfies

$$\text{Cov}_N(\hat{\pi}) \leq \text{Cov}_{N'}(\pi_{\theta^*}) + C \frac{\log \mathcal{N}_\infty(\Pi; \alpha) + \log(\delta^{-1})}{n} + \frac{Ce^{2\alpha} N'}{N} \cdot \inf_{\varepsilon > 0} \left\{ \frac{\log \mathcal{N}_\infty(\Pi, \varepsilon)}{n} + \varepsilon \right\}, \quad (33)$$

where $C > 0$ is an absolute constant.

In words, we show that for convex class, the coverage of MLE $\hat{\pi}$ can be upper bounded by the coverage of π_{θ^*} , the best-in-class approximate of π_0 . In particular, when $\pi_0 \in \Pi$, we get the following bound for convex class Π :

$$\begin{aligned} \text{Cov}_N(\hat{\pi}) &\lesssim \frac{1}{N} \cdot \inf_{\varepsilon > 0} \left\{ \frac{\log \mathcal{N}_\infty(\Pi, \varepsilon)}{n} + \varepsilon \right\} + \frac{\log \mathcal{N}_\infty(\Pi, c \log N) + \log(\delta^{-1})}{n} \\ &= \frac{\mathcal{C}_{\text{fine}}(\Pi, n)}{N^{1-2c}} + \mathcal{C}_{\text{coarse}}(\Pi, N, n), \end{aligned}$$

which improves upon the bound $\text{Cov}_N(\hat{\pi}) \lesssim \frac{\mathcal{C}_{\text{fine}}(\Pi, n)}{\log N} + \mathcal{C}_{\text{coarse}}(\Pi, N, n)$ shown in [Theorem 4.1](#) for general class Π . The proof is presented in [Appendix J.3](#).

F.3 STOCHASTIC GRADIENT DESCENT: IMPROVED GRADIENT NORMALIZATION FOR DISTILLATION

In this section, we focus on autoregressive linear models (2), and consider a variant of our setting inspired by distillation. We assume that for each example $(x^i, y_{1:H}^i)$, for each $h = 1, \dots, H$, we have access to the true next-token probabilities $\pi_D(y_h | x^i, y_{1:h-1}^i)$ for all $y_h \in \mathcal{V}$. This is an unrealistic assumption for general pre-training, but it is natural for distillation, where π_D corresponds to a teacher model (in particular, the next-token probabilities are already computed as part of a standard forward pass through the teacher model).

For the distillation setting, we give an improved gradient normalization scheme that improves upon the rate achieved by Theorem 5.1, closing the gap between SGD and maximum likelihood by matching the guarantee for Theorem 4.2.

Define $\epsilon_\theta(x, y_{1:h-1}) := D_{\text{KL}}(\pi_D(\cdot | x, y_{1:h-1}) \| \pi_\theta(\cdot | x, y_{1:h-1}))$; note that for the distillation setting, we can compute this quantity in closed form for any prefix $x, y_{1:h-1}$ in the training corpus. We consider the following (single-sample) truncated/normalized stochastic gradient estimator:

$$\hat{g}_\theta(y | x) = \sum_{h=1}^H \alpha_\theta(x, y_{1:h-1}) \nabla \log \pi_\theta(y_h | x, y_{1:h-1}), \quad (34)$$

where $A := \log N$, and where

$$\alpha_\theta(x, y_{1:h-1}) = \begin{cases} 1, & \sum_{j \leq h-1} \epsilon_\theta(x, y_{1:j}) \leq A, \\ 0, & \sum_{j < h-1} \epsilon_\theta(x, y_{1:j}) > A, \\ \frac{A - \sum_{j < h-1} \epsilon_\theta(x, y_{1:j})}{\epsilon_\theta(x, y_{1:h-1})}, & \text{otherwise.} \end{cases} \quad (35)$$

With this definition, we define the following normalized SGD update:

$$\theta^{t+1} = \text{Proj}_\Theta(\theta^t + \eta \hat{g}_{\theta^t}(y^t | x^t)). \quad (36)$$

Intuitively, the idea behind the update in Eq. (34) is to truncate the gradient at the point where the KL divergence between the teacher and student model is too large, and then normalize the gradient by the KL divergence; this is inspired by the structural result Proposition D.10 in Appendix D.4, where we show a close connection between the coverage profile and a certain “stopped” variant of KL divergence.

Theorem F.2. *Let $T, N \geq 1$ be given. With a suitably chosen stepsize $\eta > 0$, the normalized SGD update (9) achieves the following coverage bound:*

$$\mathbb{E} \left[\frac{1}{T} \sum_{t=1}^T \text{Cov}_N(\pi_{\theta^t}) \right] \lesssim \sqrt{\frac{\sigma_\star^2}{T \log N}} + \frac{B^2}{T}. \quad (37)$$

This guarantee matches the rate of Theorem 4.2 for the maximum likelihood estimator. The proof is presented in Appendix L.2.

F.4 AN IMPROVED TOURNAMENT VIA ON-POLICY GENERATION

We describe an improved tournament estimator that is able to remove that $1/N^{1-a}$ term from Theorem 6.2, meaning it achieves nontrivial guarantees even when the coverage parameter N is a constant.

Note that the term $1/N^{1-a}$ of Eq. (14) comes from the fact that $\mathbb{P}_{\pi_D}(\frac{\pi(y|x)}{\pi_D(y|x)} \geq N)$ can be as large as $1/N$ in the worst case, implying that the $\hat{\pi}$ produced by Eq. (13) may at best achieve a coverage of $1/N$. To overcome this, we introduce an *offset term*:

$$\hat{\pi} := \arg \min_{\pi \in \Pi} \max_{\pi' \in \Pi} \{ \widehat{\text{Cov}}_N(\pi' \| \pi) - 2N^a \cdot \widehat{\text{Cov}}_N(\pi' \| \pi) \}, \quad (38)$$

where we define $\widehat{\text{Cov}}_N(\pi' \| \pi) := \frac{1}{n} \sum_{i=1}^n \mathbb{P}_{y \sim \pi(\cdot | x^i)} \left(\frac{\pi'(y|x^i)}{\pi(y|x^i)} \geq N \right)$ for models $\pi, \pi', \bar{\pi}$. This estimator augments the simple tournament in Eq. (13) with an “offset” term that accounts for the fact that some of the models might be quite far from π_D . The main guarantee is as follows.

Theorem F.3. Fix $N \geq 1$, $a > 0$ such that $N^{1-2a} \geq 4$. Suppose that there exists $\bar{\pi} \in \Pi$ such that $|\log \pi_D(y | x) - \log \bar{\pi}(y | x)| \leq a \log N$ for any $x \in \mathcal{X}, y \in \mathcal{Y}$. Then with probability $1 - \delta$, the tournament estimator (38) achieves $\text{Cov}_{2N^{1+a}}(\hat{\pi}) \lesssim \frac{\log(|\Pi|/\delta)}{n}$.

Compared to Theorem 6.2, this tournament eliminates the additive $1/N^{1-a}$ term. It does, however, require a stronger condition on the best-in-class model $\bar{\pi}$ that $|\log \pi_D(y | x) - \log \bar{\pi}(y | x)| \leq a \log N$, which implies in particular that $\text{Cov}_{N^a}(\bar{\pi}) = 0$.

Infinite classes: Beating maximum likelihood. While we motivated the tournament estimators through model/checkpoint selection with a finite class Π , both estimators can also be applied to general, infinite classes Π . In this case, it turns out that they both improve upon the coverage achieved by the maximum likelihood estimator in Theorem 4.1, even in the well-specified case where $\pi_D \in \Pi$; informally, the tournament estimators allow us to remove the fine-grained term in Theorem 4.1, leaving only a coarse-grained term. See Theorem 6.2' and Theorem F.3' for the formal statements.

G EXPERIMENTS

We describe the general graph search task used throughout our experiments in [Appendix G.1](#), then detail the specific setups used for [Figure 1](#) in [Appendix G.2](#), and for [Figure 2](#) in [Appendix G.3](#).

G.1 GRAPH REASONING TASK

We evaluate our theoretical predictions using experiments in graph reasoning tasks, in which transformer models are trained to find paths between source and target nodes in graphs. Both graph reasoning benchmarks and synthetic datasets have seen increasing use as abstractions for reasoning problems and for probing language modeling phenomena ([Sanford et al., 2024](#); [Nagarajan et al., 2025](#); [Saparov et al., 2025](#); [Bachmann & Nagarajan, 2024](#); [Yehudai et al., 2025](#); [Taylor et al., 2024](#); [Wang et al., 2023](#); [Fatemi et al., 2024](#); [Tang et al., 2025](#)).

These tasks provide minimal abstractions of core reasoning problems, yet are expressive enough to capture pre-training and fine-tuning phenomena. They also offer flexibility in problem structure and difficulty: by specifying different graph topologies and path depths, we can modulate difficulty and expose sources of hardness. At the same time, the simplicity of the setup enables training in controlled settings with interpretable results with which to ground our theoretical predictions.

G.1.1 GRAPH SEARCH TASK DESCRIPTION

The graph search tasks in [Appendix G.2](#) and [Appendix G.3](#) share the same high-level structure, and are comprised of

- A set of graph structures $\mathcal{G}$ that map bijectively to a set of prompts $\mathcal{X}$, and induce the response space $\mathcal{Y}$
- A distribution over the prompts $\mu \in \Delta(\mathcal{X})$
- A data collection policy $\pi_D : \mathcal{X} \rightarrow \Delta(\mathcal{Y})$

Next, we describe the general details of the graph search task common to all experiments, and leave the task details for each figure in the proceeding sections.

Graph Search. The nodes of all graphs in a given task $\mathcal{G}$ are drawn from the set $[m]$, for some integer $m \in \mathbb{N}$. Each graph structure $G = (V, E) \in \mathcal{G}$ is comprised of a set of vertices (“nodes”) $V \subseteq [m]$ and edges $E = \{(u, v) : u, v \in V, u \neq v\}$. It also contains one source node $s \in V$ and one target node $t \in V$, so that (G, s, t) specifies one search problem instance within the class. The search task can be translated into an autoregressive sequence modeling problem using the below prompt and response specifications.

Layered Graph Structure. For all experiments, we utilize a *layered directed acyclic graph (layered DAG)* with a rectangular structure. Following the source node, the graph has L layers each with a fixed number of nodes. In each layer, only a subset of its nodes has edges connecting to the next layer, and we refer to these nodes as *passable nodes*, or the set $\{v \in V : \deg^+(v) > 0\}$ that has non-zero out-degree. In a given layer each passable node has edges to all nodes in the next layer, while the remaining (non-passable) nodes in the layer cannot be used to traverse to the target, so that as soon as the model outputs one such node its path cannot be valid.

In order to output a valid path from source to target, it is sufficient to keep outputting the passable nodes in the next layer. In general a graph may have many valid paths, and in each experiment, π_D always samples valid paths. However, as will describe shortly, π_D may use complex functions to sample from only a subset of the valid paths, and $\hat{\pi}$ must learn to cover such behavior.

The layered DAG offers a natural interpretation as an abstraction for reasoning problems. Following passable nodes in valid paths corresponds to taking reasoning steps that make progress towards the solution, and selecting paths via more complex functions maps to learning high-quality solutions that accurately reflect desired properties for the problem’s solution.

Graphs to Prompts. Given a graph structure $G = (V, E) \in \mathcal{G}$ and a source node $s \in V$ and target node $t \in V$, the prompt x encodes the search problem as the adjacency list of G with the source and target nodes appended to the end. The prompt is formatted as a string of the form

$$x : u_1 \ v_1 \mid u_2 \ v_2 \mid \dots \mid u_k \ v_k \ / \ s \ t =$$

where $(u_i, v_i) \in [m]^2$ are the vertices of the i -th edge in the adjacency list. For example, with edges $E = \{(10, 23), (86, 47), \dots, (45, 32)\}$, the prompt is formatted as

$$x: 10\ 23\ |\ 86\ 47\ |\ \dots\ |\ 45\ 32\ /\ 10\ 45\ =$$

where $|$ and $/$ are special characters that separate two edges and the adjacency list from the source and target nodes, respectively. The special character $=$ indicates the end of the prompt.

Graphs to Responses. Given a graph structure $G = (V, E) \in \mathcal{G}$ and a source node $s \in V$ and target node $t \in V$, the response y encodes the path from the source to the target node in G . In general a graph may have multiple paths from source to target. The horizon H corresponds to the longest path length in $\mathcal{G}$, and a response takes the form of a string

$$y: s\ v_{-1}\ v_{-2}\ v_{-3}\ \dots\ v_{-H}\ t$$

where $v_i \in [m]$ are the vertices of the i -th edge in the path.

Sequence Modeling Problem. In summary, a graph search task with set of graphs $\mathcal{G}$ induces an autoregressive sequence modeling problem with a vocabulary space $\mathcal{V} = [m] \cup \{ |, /, = \}$, prompts $\mathcal{X} \subseteq \mathcal{V}^*$ corresponding to graph structures, and responses $\mathcal{Y} \subseteq \mathcal{V}^H$ corresponding to paths with length at most H . In addition, the task is equipped with $\mu \in \Delta(\mathcal{X})$ and $\pi_D: \mathcal{X} \rightarrow \Delta(\mathcal{Y})$ that is used to collect the training dataset $D = \{(x, y)\}$, where $x \sim \mu$ and $y \sim \pi_D(x)$.

G.1.2 GENERAL IMPLEMENTATION DETAILS

Next, we describe the common implementation details of the graph search task.

Tokenizer. The tokenizer is a numeral tokenizer standard for graph reasoning tasks. Each node $v \in [m]$ is tokenized as its integer node value, and the special characters $|$, $/$, and $=$ are tokenized as $m+1, m+2, m+3$, respectively.

Transformer model. Throughout our experiments, we train causally-masked GPT2 transformer models to minimize the cross-entropy loss using the Adam optimizer with fixed learning rate, and perform a grid search over the parameters displayed in Table 1. Parameters with fixed values were chosen based on related papers such as Bachmann & Nagarajan (2024). In both experiments, the model architecture with 4 heads, 6 hidden layers, and 384 hidden dimensions worked best. We use absolute positional encodings. Training iterations and grid search values for the learning rate are different for each experiment, and discussed further below.

Hyperparameter	Values
Number of heads	{4, 6, 8}
Number of layers	{3, 4, 6, 8}
Hidden dimensions	384
Activation function	GeLU
Batch size	128
Weight decay	0.01

Table 1: Hyperparameter grid search values for transformer models in graph search.

G.2 EXPERIMENT DETAILS FOR FIGURE 1

The graph search task for Figure 1 exposes natural properties of pre-training data where cross-entropy reduction comes at the cost of a worse coverage profile. In particular, because pre-training data is diverse, the model in practice is generally unable to perfectly fit the distribution. When one mode of behavior is better-represented than another, cross-entropy minimization, which is an average-case distribution-matching metric, can sacrifice coverage over the different modes in order to increase performance on one.

Correspondingly, our graph search task for Figure 1 is a mixture of two classes of graph structures. Due to representational and finite-sample constraints, the model is unable to fit both perfectly during training, and, in particular, fitting one class well (in the sense of cross-entropy loss) comes at the cost of worse performance on the other. The checkpoint with the best coverage arises at some middle point in training when the model learns both classes of graphs equally well, and has good coverage over

both classes (the dip Cov_N in the leftmost subplot of Figure 1). Further reduction of cross-entropy loss over the latter half of training requires the model to lose coverage over π_D in the less-represented graph class (observed as the increase in Cov_N in the latter half of training iterations).

Even though the task cannot be learned perfectly from the supervised learning feedback, the model can still learn a policy that always samples a correct path matching π_D ’s with $O(1)$ sampling attempts, which means that it leads to efficient downstream post-training (e.g., on one of modes or with reward-based feedback), and also achieve optimal performance with test-time inference methods.

For the experiments in Figure 1, we first pre-train a model on a larger set of graph structure classes so that it learns a diverse set of behaviors, then finetune its behavior on two. The performance on the finetuning task is displayed in Figure 1, and we first describe the finetuning dataset, followed by the pre-training dataset.

G.2.1 TASK DESCRIPTION

The finetuning task is a skewed mixture over two graph disjoint graph classes, $\mathcal{G}_1 \cup \mathcal{G}_2 = \mathcal{G}$, with $\tilde{\mu} \in \Delta(\{1, 2\})$ denoting the probability of each class in the data. All graphs in $\mathcal{G}$ follow the *layered DAG* structure described in Appendix G.1, with $L = 8$ layers that each have 4 nodes. Of the 8 layers, 2 are randomly selected to have 2 passable nodes (meaning that they are connected to the next layer), while the remaining layers have only 1 passable node. Although there are 4 valid paths from source to target, the policy π_D is deterministic and chooses 1 based on a rule, which is what distinguishes the two class types described below.

Class $\mathcal{G}_1$ with probability $\tilde{\mu}(1) = 0.9$. For an integer $j \in \mathbb{Z}$, let the function $p(j) = (j \bmod 2)$ denote its parity. In layers with 1 passable node, π_D takes the passable node. For each layer $l \in [L]$ with 2 passable nodes (there is guaranteed to be one even and one odd), π_D chooses the node v such that $p(v) = p(l)$, that is, the node whose parity is equal to the parity of the layer index.

Class $\mathcal{G}_2$ with probability $\tilde{\mu}(2) = 0.1$. In this class π_D takes the opposite rule from the one in $\mathcal{G}_2$: for layers l with 2 passable nodes, it chooses the node v such that $p(v) = 1 \oplus p(l)$.

The class of a graph is technically identifiable from the prompt, but the problem is too difficult for the model to learn in just the finetuning stge. The class of a graph can be computed from the parity of a hidden subset of their nodes whose cardinality is half the total number of nodes; letting this hidden subset be $V' \in V$, all graphs in $\mathcal{G}_1$ have $1 = \bigoplus_{u \in V'} p(u)$, while the opposite is true for all graphs in $\mathcal{G}_2$.

Dataset generation. Each sample in the dataset $D = \{(x, y)\}$ is then generated via the following procedure.

1. First sample an index $i \sim \tilde{\mu}$.
2. Sample $G \in \mathcal{G}_i$ by randomly drawing $V \subset [m]$ without replacement, and instantiate the edges according to the description for each class above.
3. Format the prompt x per Appendix G.1.
4. Draw $y \sim \pi_D(\cdot \mid x)$ according to description for each class above.

G.2.2 PRE-TRAINING DESCRIPTION

The graphs in the pre-training task are a superset of the graphs in the finetuning task, that is, $\cup_{i \in [K]} \mathcal{G}_i = \mathcal{G}$ with $K = 3$, and the data distribution is a uniform mixture of these 3 classes, $\tilde{\mu}(i) = \frac{1}{K}$ for each $i \in [K]$. The first two classes $\mathcal{G}_1$ and $\mathcal{G}_2$ are defined exactly as they are in the finetuning dataset. In $\mathcal{G}_3$, two layers have 2 passable nodes, while the rest have 1, and π_D samples one of the 2^2 valid paths from source to target at random. The dataset is sampled using the same dataset generation procedure described for the finetuning task above.

G.2.3 TASK-SPECIFIC IMPLEMENTATION DETAILS

The transformer model is first pre-trained on a fixed dataset drawn from the pre-training distribution, with $8 \times 64,000$ prompts in total, using a learning rate of $1e-4$ for 200k iterations, which was chosen based on a grid search over learning rates $\{5e-5, 1e-4, 5e-4\}$.

The final checkpoint is then finetuned for 50k iterations in an online fashion, where fresh samples are drawn for each batch (this is equivalent to offline training with a dataset tha has an equivalent number

of samples). The learning rate is $5e-6$, which was chosen based on a grid search over learning rates $\{5e-6, 1e-5\}$.

G.3 EXPERIMENT DETAILS FOR [FIGURE 2](#)

The graph structures used for [Figure 2](#) to expose the dependence on horizon of KL-divergence but not Cov_N leverages the intuition from [Remark 3.1](#). The training data is homogeneous, and a fraction consists of difficult graph problems that the learner cannot cover with the given finite samples. The coverage profile Cov_N will be the same regardless of H , but KL scales linearly with H due to this unlearnable subset of the data.

G.3.1 TASK DESCRIPTION

For [Figure 2](#), we devise a family of tasks that is defined per H , which we then instantiate with $H \in \{8, 16, 24\}$ for our results. For a fixed H , the task $\mathcal{G}_H$ utilizes the *layered DAG* graph structure described in [Appendix G.1](#) with H layers of 4 nodes each, so that $\mathcal{Y} = \mathcal{V}^{H+2}$ when the source and target nodes are included.

The task is a heterogeneous mixture over 3 classes of graphs described below that we refer to as $\mathcal{G}_{H,1} \cup \mathcal{G}_{H,2} \cup \mathcal{G}_{H,3} = \mathcal{G}_H$. The classes $\mathcal{G}_{H,2}$ and $\mathcal{G}_{H,3}$ are significantly harder to learn and the model will fail to do so with the given number of training samples, even though $\mathcal{G}_{H,1}$ is learned quickly (and also provides useful features for learning the other two tasks). The distribution over these 3 classes is fixed for all H and specified by $\tilde{\mu} \in \Delta(\{1, 2, 3\})$.

Class $\mathcal{G}_{H,1}$ with probability $\tilde{\mu}(1) = 0.94$. All H layers have only 1 passable node, so each $G \in \mathcal{G}_{H,1}$ has only one valid path from source to target. For prompts corresponding to graphs in this class, π_D deterministically takes the single valid path.

Class $\mathcal{G}_{H,2}$ with probability $\tilde{\mu}(2) = 0.05$. Half of the layers (or $H/2$, selected randomly) have 2 passable nodes while the rest have 1. While there are $2^{H/2}$ valid paths from source to target, π_D deterministically selects one of them. For layers with 2 passable nodes, one is guaranteed to be even and the other odd. In layers with more than one passable node, π_D selects one node by following a difficult, deterministic rule. This rule requires π_D to select the node v whose parity matches the parity of the layer index, XOR’ed with the parity of each passable node in the entire graph. More specifically, recall that $p(j)$ denotes the parity of an integer $j \in [m]$, and let $V^* := \{v \in V : \deg^+(v) > 0\} \subset V$ denote the set of all passable nodes (or those with positive out-degree). Then in layer l , π_D selects the node v such that $p(v) = p(l) \oplus (\bigoplus_{u \in V^*} p(u))$.

Class $\mathcal{G}_{H,3}$ with probability $\tilde{\mu}(3) = 0.01$. Regardless of H , 4 of the layers are randomly chosen to have 2 passable nodes, so that there are $2^4 = 16$ valid paths from source to goal. Here, however, π_D samples one of the $2^{H/2}$ valid paths uniformly at random.

Note that prompts/graphs from each class are distinguishable from each other (or, identifiable) based on prompt features alone, so a powerful-enough model can achieve perfect performance across all of them simultaneously. $\mathcal{G}_{H,2}$, for example, has more edges and thus a longer prompt than $\mathcal{G}_{H,1}$; similar statements apply to $\mathcal{G}_{H,3}$. Dataset generation occurs in the same manner as described in [Appendix G.2](#).

G.3.2 TASK-SPECIFIC IMPLEMENTATION DETAILS

Here, we describe experiment-specific implementation details on top of those previously described in [Appendix G.1](#) (which apply to all figures).

In addition to a grid search over the parameters in [Table 1](#), we perform a search over learning rates $\{5e-5, 1e-4, 5e-4\}$, for which the learning rate of $1e-4$ exhibited the best validation performance. The model is trained for 40k iterations over a fixed dataset of $8 \times 64,000$ samples.

The results in [Figure 2](#) are computed from evaluations of training checkpoints on per-class validation datasets of 1024 prompts from each $\mathcal{G}_{H_i}$; these metrics are then averaged according to the probabilities in $\tilde{\mu}$ to obtain the final result. In total we ran 16 seeds, and plot their median. The shaded region in [Figure 2](#) displays the region between the $\frac{1}{16}$ quantile and $\frac{15}{16}$ quantile.

Part II

Proofs

H TECHNICAL TOOLS

H.1 CONCENTRATION INEQUALITIES

Lemma H.1 (Azuma-Hoeffding). *Let $(Z^i)_{i \leq n}$ be a sequence of real-valued random variables adapted to a filtration $(\mathcal{F}_i)_{i \leq n}$. If $|Z^i| \leq R$ almost surely, then with probability at least $1 - \delta$, for all $n' \leq n$,*

$$\left| \sum_{i=1}^{n'} Z^i - \mathbb{E}_{i-1}[Z^i] \right| \leq R \cdot \sqrt{8n \log(2\delta^{-1})}.$$

Lemma H.2 (Freedman's inequality). *Let $(Z^i)_{i \leq n}$ be a real-valued martingale difference sequence adapted to a filtration $(\mathcal{F}_i)_{i \leq n}$. If $|Z^i| \leq R$ almost surely, then for any $\eta \in (0, 1/R)$, with probability at least $1 - \delta$, for all $n' \leq n$,*

$$\sum_{i=1}^{n'} Z^i \leq \eta \sum_{i=1}^{n'} \mathbb{E}_{i-1}[(Z^i)^2] + \frac{\log(\delta^{-1})}{\eta}.$$

The next result is a standard consequence of [Lemma H.2](#) (e.g., [Foster et al. \(2021\)](#)).

Lemma H.3. *Let $(Z^i)_{i \leq n}$ be a sequence of random variables adapted to a filtration $(\mathcal{F}_i)_{i \leq n}$. If $0 \leq Z^i \leq R$ almost surely, then with probability at least $1 - \delta$, for all $n' \leq n$,*

$$\sum_{i=1}^{n'} Z^i \leq \frac{3}{2} \sum_{i=1}^{n'} \mathbb{E}_{i-1}[Z^i] + 4R \log(2\delta^{-1}), \quad (39)$$

and

$$\sum_{i=1}^{n'} \mathbb{E}_{i-1}[Z^i] \leq 2 \sum_{i=1}^{n'} Z^i + 8R \log(2\delta^{-1}). \quad (40)$$

Lemma H.4. *Suppose that μ is a distribution over $\mathcal{Z}$, function class $\mathcal{F} \subseteq (\mathcal{Z} \rightarrow \mathbb{R})$ is given. We let $N(\mathcal{F}, \epsilon; \|\cdot\|_\infty)$ be the ϵ -covering number of $\mathcal{F}$ under the norm $\rho(f, f') := \sup_{z \in \mathcal{Z}} |f(z) - f'(z)|$. Then, under $\mathcal{D} = \{Z^1, \dots, Z^n\}$ drawn from μ i.i.d, the following holds with probability at least $1 - \delta$:*

$$\sum_{i=1}^n f(Z^i) \leq n \log \mathbb{E}_\mu[\exp(f(Z))] + \inf_{\epsilon \geq 0} \{\log N(\mathcal{F}, \epsilon; \|\cdot\|_\infty) + 2Ln\epsilon\}.$$

Lemma H.5. *For distribution $P, Q \in \Delta(\mathcal{X})$, function $f : \mathcal{X} \rightarrow [-B, B]$, it holds that*

$$\|\mathbb{E}_P[f] - \mathbb{E}_Q[f]\|^2 \leq 3\text{Var}_Q[f] \cdot D_H^2(P, Q) + 8B^2 D_H(P, Q)^4.$$

Therefore, for any $f : \mathcal{X} \rightarrow \mathbb{R}^d$ with $\|f\| \leq B$, it holds that

$$\|\mathbb{E}_P[f] - \mathbb{E}_Q[f]\| \leq 2\sqrt{\mathbb{E}_Q\|f - \mathbb{E}_Q[f]\|^2} \cdot D_H(P, Q) + 3BD_H^2(P, Q), \quad (41)$$

and

$$\mathbb{E}_P\|f - \mathbb{E}_P[f]\|^2 \leq 3\mathbb{E}_Q\|f - \mathbb{E}_Q[f]\|^2 + 8B^2 D_H^2(P, Q). \quad (42)$$

Proof of Lemma H.5. We denote $P(x)$ ($Q(x)$) to be the density function of P (Q). Then

$$|\mathbb{E}_P[f] - \mathbb{E}_Q[f]|^2 = \left(\int_{\mathcal{X}} (f(x) - \mathbb{E}_Q[f])(P(x) - Q(x))dx \right)^2 \quad (43)$$

$$\leq \int_{\mathcal{X}} (f(x) - \mathbb{E}_Q[f])^2 (\sqrt{P(x)} + \sqrt{Q(x)})^2 dx \cdot \int_{\mathcal{X}} (\sqrt{P(x)} - \sqrt{Q(x)})^2 dx \quad (44)$$

$$\leq 4D_H^2(P, Q) \cdot \left(\text{Var}_Q[f] + \mathbb{E}_P(f - \mathbb{E}_Q[f])^2 \right) \quad (45)$$

Further, we know

$$\mathbb{E}_P(f - \mathbb{E}_Q[f])^2 \leq 3 \mathbb{E}_Q(f - \mathbb{E}_Q[f])^2 + 8B^2 D_H^2(P, Q). \quad (46)$$

This gives the desired upper bound. $\square$

Lemma H.6. Suppose that $\phi : \mathcal{Y} \rightarrow \mathbb{B}_2(B)$ with $B \geq 1$, and for any $\theta \in \mathbb{B}_2(1)$, $\pi_\theta \in \Delta(\mathcal{Y})$ is defined as $\pi_\theta(y) \propto \exp(\langle \phi(y), \theta \rangle)$. Then for any $\theta^*, \theta \in \mathbb{B}_2(1)$, it holds that

$$\mathbb{E}_{y \sim \pi_{\theta^*}} \langle \phi(y) - \mathbb{E}_{\pi_{\theta^*}}[\phi], \theta - \theta^* \rangle^2 \leq 15B D_{\text{KL}}(\pi_{\theta^*} \parallel \pi_\theta).$$

Proof. Denote $\bar{\phi}(y) := \phi(y) - \mathbb{E}_{\pi_{\theta^*}}[\phi]$. By definition,

$$D_{\text{KL}}(\pi_{\theta^*} \parallel \pi_\theta) = \log \mathbb{E}_{y \sim \pi_{\theta^*}} [\exp(\langle \bar{\phi}(y), \theta - \theta^* \rangle)] \geq B \log \mathbb{E}_{y \sim \pi_{\theta^*}} \left[\exp\left(\frac{1}{B} \langle \bar{\phi}(y), \theta - \theta^* \rangle\right) \right]$$

Note that for $x \geq -4$, we have $e^x \geq 1 + x + \frac{1}{10}x^2$. Therefore, we have

$$\frac{1}{B} D_{\text{KL}}(\pi_{\theta^*} \parallel \pi_\theta) \geq \log \left(1 + \frac{1}{10B^2} \mathbb{E}_{y \sim \pi_{\theta^*}} \langle \bar{\phi}(y), \theta - \theta^* \rangle^2 \right) \geq \frac{1}{15B^2} \mathbb{E}_{y \sim \pi_{\theta^*}} \langle \bar{\phi}(y), \theta - \theta^* \rangle^2,$$

where we use $\log(1+x) \geq \frac{3}{4}x$ for all $x \in [0, \frac{8}{5}]$. $\square$

I PROOFS FROM SECTION 3

Proof of Proposition 3.2. Consider the setting $d = 1$, $\mathcal{X} = \{0, 1\}$, $\mathcal{Y} = \{-B, B\}$, the distribution μ be given by $\mu(1) = 1 - \mu(0) = \frac{1}{2n}$, and the feature map $\phi : \mathcal{X} \times \mathcal{Y}^* \rightarrow [-B, B]$ be given by $\phi(0, \cdot) = 0$, and $\phi(1, y_{1:h}) = y_h$.

Note that under this construction, $\mathbb{P}_{\{x^t\}_{t \in [n]} \sim \pi_0} (x^t = 0 \forall t \in [T]) \geq 1 - n\mu(1) = \frac{1}{2}$. We let E be the event $\{x^t = 0 \forall t \in [T]\}$. Then, for any $\theta^* \in [-1, 1]$,

$$\mathbb{E}_{\mathcal{D} \sim \pi_{\theta^*}} [D_{\text{KL}}(\pi_{\theta^*} \parallel \hat{\pi}) \mid E] = \mathbb{E}_{\mathcal{D} \sim \pi_0} [D_{\text{KL}}(\pi_{\theta^*} \parallel \hat{\pi}) \mid E].$$

Furthermore, for any $\hat{\pi} \in \Pi$,

$$D_{\text{KL}}(\pi_{\theta^*} \parallel \hat{\pi}) = H \cdot \mu(1) \cdot D_{\text{KL}}(\pi_{\theta^*}(y_1 = \cdot \mid x = 1) \parallel \hat{\pi}(y_1 = \cdot \mid x = 1)),$$

and hence,

$$D_{\text{KL}}(\pi_1 \parallel \hat{\pi}) + D_{\text{KL}}(\pi_{-1} \parallel \hat{\pi}) \geq \frac{H}{2n} \cdot 2D_{\text{KL}}\left(\text{Ber}\left(\frac{e^B}{e^B + e^{-B}}\right) \parallel \text{Ber}\left(\frac{1}{2}\right)\right) \geq \frac{H}{2n}.$$

Therefore, we can lower bound

$$\begin{aligned} & \mathbb{E}_{\mathcal{D} \sim \pi_1} [D_{\text{KL}}(\pi_1 \parallel \hat{\pi})] + \mathbb{E}_{\mathcal{D} \sim \pi_{-1}} [D_{\text{KL}}(\pi_{-1} \parallel \hat{\pi})] \\ & \geq \mathbb{P}(E) \cdot \mathbb{E}_{\mathcal{D} \sim \pi_0} [D_{\text{KL}}(\pi_1 \parallel \hat{\pi}) + D_{\text{KL}}(\pi_{-1} \parallel \hat{\pi}) \mid E] \geq \frac{1}{2} \cdot \frac{H}{2n}. \end{aligned}$$

This gives the desired lower bound. $\square$

Note that in the construction above, the variance $\sigma_\star^2$ (defined in Section 4.1) can be bounded by $\sigma_\star^2 \lesssim \frac{He^{-2B}}{n}$.

J PROOFS FROM SECTION 4

J.1 PROOF SKETCH FOR THEOREM 4.1

Fix $N \geq 8$ and let $\varepsilon \in [0, 1]$ be a parameter to be set later. Let $\Pi_{\text{bad}}(\varepsilon) := \{\pi \in \Pi \mid \text{Cov}_N(\pi) \geq \varepsilon\}$ be the set of $\pi \in \Pi$ that fail to achieve coverage ε . The basic idea behind the proof of Theorem 4.1 is to interpret the condition $\text{Cov}_N(\pi) \geq \varepsilon$ as an small-ball like *anti-concentration* condition in the vein of Mendelson (2014; 2017). That is, for models $\pi \in \Pi_{\text{bad}}(\varepsilon)$ where coverage fails, the condition $\text{Cov}_N(\pi) \geq \varepsilon$ witnesses a *one-sided* tail bound which implies that the empirical likelihood of π is *not too large* with high probability, which means that $\pi \in \Pi_{\text{bad}}(\varepsilon)$ cannot be a maximum-likelihood solution.

Let $\mathcal{S}_N(\pi) := \frac{1}{n} |\{i \in [n] \mid \frac{\pi_D(y^i | x^i)}{\pi(y^i | x^i)} \geq N^{1-2c}\}|$ denote the empirical probability of π fails to cover π_D . Our first step is to show via covering and concentration that with high-probability, all $\pi \in \Pi$ satisfy

$$\mathcal{S}_N(\pi) \geq \frac{1}{2} \text{Cov}_N(\pi) - \mathcal{C}_{\text{coarse}}(\Pi, N, n), \quad (47)$$

Here we only pays the covering number at a coarse scale (leading to the coarse-grained term in Theorem 4.1) because we only need to show that coverage concentrates, not the log-loss itself.

Eq. (47) implies that for all $\pi \in \Pi_{\text{bad}}(\varepsilon)$, we can bound

$$\begin{aligned} \hat{L}_n(\pi) - \hat{L}_n(\pi_D) &= - \sum_{i=1}^n \left[\log \frac{\pi_D(y^i | x^i)}{\pi(y^i | x^i)} - C \right]_+ + \sum_{i=1}^n \log \frac{\pi(y^i | x^i)}{\pi_D(y^i | x^i)} \vee C \\ &\leq -|\mathcal{S}_N(\pi)|((1-2c) \log N - C) + \sum_{i=1}^n \log \frac{\pi(y^i | x^i)}{\pi_D(y^i | x^i)} \vee C \\ &\leq -\frac{n}{4} \log N \cdot \text{Cov}_N(\pi) + \sum_{i=1}^n \log \frac{\pi(y^i | x^i)}{\pi_D(y^i | x^i)} \vee C, \end{aligned}$$

where $C > 0$ is any fixed constant. Finally, using a variation of a standard one-sided tail bound for the logarithmic loss (van de Geer, 2000; Zhang, 2006),⁷ we show that with high probability, all $\pi \in \Pi$ satisfy

$$\sum_{i=1}^n \log \frac{\pi(y^i | x^i)}{\pi_D(y^i | x^i)} \vee C \leq \mathcal{C}_{\text{fine}}(\Pi, n) \cdot n,$$

as long as $C \geq \log 4$. Combining these results, we conclude that

$$\text{Cov}_N(\pi) \lesssim \frac{\hat{L}_n(\pi_D) - \hat{L}_n(\pi) + \mathcal{C}_{\text{fine}}(\Pi, n)}{n \log N} + \mathcal{C}_{\text{coarse}}(\Pi, N, n).$$

It follows that if $\varepsilon \gtrsim \frac{1}{\log N} \cdot \mathcal{C}_{\text{fine}}(\Pi, n) + \mathcal{C}_{\text{coarse}}(\Pi, N, n)$, then all $\pi \in \Pi_{\text{bad}}(\varepsilon)$ have $L(\pi) - L(\pi_D) < 0$, and since $\pi_D \in \Pi$, this means that no such $\pi \in \Pi_{\text{bad}}(\varepsilon)$ can be the maximum likelihood solution.

J.2 PROOF OF THEOREM 4.1

Theorem 4.1' (General version of Theorem 4.1). *Let $N \geq 8$ be given. With probability at least $1 - \delta$, any approximate maximum likelihood estimator $\hat{\pi}$ with $\hat{L}_n(\hat{\pi}) \geq \max_{\pi \in \Pi} \hat{L}_n(\pi) - n\varepsilon_{\text{apx}}$ satisfies*

$$\text{Cov}_N(\hat{\pi}) \lesssim \frac{\log \mathcal{N}_{\infty}(\Pi; c \log N) + \log(\delta^{-1})}{n} + \frac{1}{\log N} \left(\inf_{\varepsilon > 0} \left\{ \frac{\log \mathcal{N}_{\infty}(\Pi, \varepsilon)}{n} + \varepsilon \right\} + \varepsilon_{\text{apx}} \right), \quad (48)$$

where $c > 0$ is an absolute constant.

⁷That the bound is one-sided is critical, as this allows us to avoid paying for the range of the density ratios under consideration. For details, see Proposition J.1.

In the following, for a fixed threshold $C \geq 4$, we define the clipped log loss as

$$L_C^+(\pi) := \sum_{i=1}^n \max \left\{ \log \frac{\pi(y^i | x^i)}{\pi_D(y^i | x^i)}, -\log C \right\}, \quad (49)$$

$$L_C^-(\pi) := \sum_{i=1}^n \max \left\{ 0, \log \frac{\pi_D(y^i | x^i)}{\pi(y^i | x^i)} - \log C \right\}. \quad (50)$$

Note that $\hat{L}_n(\pi) = L_C^+(\pi) - L_C^-(\pi)$, and hence for approximate maximum likelihood estimator $\hat{\pi}$ with $\hat{L}_n(\hat{\pi}) \geq \max_{\pi \in \Pi} \hat{L}_n(\pi) - n\varepsilon_{\text{apx}}$, we have

$$L_C^-(\hat{\pi}) \leq L_C^+(\hat{\pi}) + n\varepsilon_{\text{apx}}.$$

In the following, we argue that $L_C^-(\pi)$ upper bound the coverage $\text{Cov}_N(\pi)$ for any $\pi \in \Pi$ and $M > C$, and $L_C^+(\pi)$ can be bounded by the uniform convergence argument.

Proposition J.1. *Suppose that $C \geq 4$. Then, with probability at least $1 - \delta$, it holds that for any $\pi \in \Pi$,*

$$L_C^+(\pi) \leq 2 \inf_{\epsilon \geq 0} \{ \log \mathcal{N}_\infty(\Pi, \epsilon) + n\epsilon \}.$$

Proposition J.2. *Fix any $\alpha \in (0, \frac{\log(N/C)}{2})$. Then, with probability at least $1 - \delta$, it holds that*

$$\text{Cov}_N(\pi) \leq \frac{2}{\log(N/C) - 2\alpha} \cdot L_C^-(\pi) + 2 \log(\mathcal{N}_\infty(\Pi, \alpha)/\delta).$$

The proof of [Theorem 4.1](#) and [Theorem 4.1'](#) is hence completed by combining the propositions above and setting $\alpha = c \log N$. $\square$

Proof of Proposition J.1. This is a direct corollary of [Lemma H.4](#). For each $\pi \in \Pi$, we let $f_\pi(x, y) := \max \left\{ \log \frac{\pi(y|x)}{\pi_D(y|x)}, -\log C \right\}$, and then $N(\mathcal{F}, \epsilon; \|\cdot\|_\infty) \leq \mathcal{N}_\infty(\Pi, \epsilon)$ for any $\epsilon \geq 0$. Applying [Lemma H.4](#) with [Lemma J.1](#) gives the desired upper bound. $\square$

Lemma J.1. *As long as $C \geq 4$, it holds that*

$$\mathbb{E}_{(x,y) \sim \pi_D} \exp \left(\frac{1}{2} \max \left\{ \log \frac{\pi(y|x)}{\pi_D(y|x)}, -\log C \right\} \right) \leq 1. \quad (51)$$

Proof of Lemma J.1. We denote $E := \left\{ (x, y) : \frac{\pi(y|x)}{\pi_D(y|x)} \geq \frac{1}{C} \right\}$. Then it holds that

$$\begin{aligned} & \mathbb{E}_{(x,y) \sim \pi_D} \exp \left(\frac{1}{2} \max \left\{ \log \frac{\pi(y|x)}{\pi_D(y|x)}, -\log C \right\} \right) \\ &= \mathbb{E}_{(x,y) \sim \pi_D} \left[\sqrt{\frac{\pi(y|x)}{\pi_D(y|x)}} \mathbb{I}\{(x, y) \in E\} + \frac{1}{\sqrt{C}} \mathbb{I}\{(x, y) \notin E\} \right] \\ &= \mathbb{E}_{x \sim \pi_D} \left[\sum_{y: (x,y) \in E} \sqrt{\pi(y|x) \pi_D(y|x)} \right] + \frac{1}{\sqrt{C}} \pi_D(E^c) \end{aligned}$$

By Cauchy inequality, we have

$$\sum_{y: (x,y) \in E} \sqrt{\pi(y|x) \pi_D(y|x)} \leq \sqrt{\sum_{y: (x,y) \in E} \pi(y|x) \cdot \sum_{y: (x,y) \in E} \pi_D(y|x)} \leq \sqrt{\pi_D(E)}.$$

Therefore, as long as $C \geq 4$, it holds that

$$\mathbb{E}_{(x,y) \sim \pi_D} \exp \left(\frac{1}{2} \max \left\{ \log \frac{\pi_D(y|x)}{\pi_D(y|x)}, -\log C \right\} \right) \leq \sqrt{\pi_D(E)} + \frac{1}{2}(1 - \pi_D(E)) \leq 1,$$

where we use $1 - p = (1 + \sqrt{p})(1 - \sqrt{p}) \leq 2(1 - \sqrt{p})$ for any $p \in [0, 1]$. $\square$

Proof of Proposition J.2. Fix any $N \geq 1, \alpha \geq 0$. By definition, for any $\pi \in \Pi$,

$$\begin{aligned} L_C^-(\pi) &= \sum_{i=1}^n \max \left\{ 0, \log \frac{\pi_D(y^i | x^i)}{\pi(y^i | x^i)} - \log C \right\} \\ &\geq (\log N - \log C) \left| \left\{ i \in [n] : \log \frac{\pi_D(y^i | x^i)}{\pi(y^i | x^i)} \geq \log N \right\} \right| \\ &=: n(\log N - \log C) \cdot \widehat{\text{Cov}}_M(\pi_D \| \pi), \end{aligned}$$

where we denote (cf. Lemma J.2)

$$\widehat{\text{Cov}}_N(\pi_D \| \pi) = \frac{1}{n} \left| \left\{ t \in [n] : \frac{\pi_D(y^t | x^t)}{\pi(y^t | x^t)} \geq N \right\} \right|.$$

Then, by Lemma J.2, it holds that with probability at least $1 - \delta$, for any $\pi \in \Pi$,

$$\widehat{\text{Cov}}_N(\pi_D \| \pi) \geq \frac{1}{2} \text{Cov}_{e^{2\alpha} N}^{\pi_D}(\pi_D \| \pi) - \log(\mathcal{N}_\infty(\Pi, \alpha)/\delta).$$

Rescaling $N \leftarrow e^{-2\alpha} N$ and reorganizing complete the proof. $\square$

Lemma J.2. For any policy π, π' , we consider the quantities

$$\widehat{\text{Cov}}_N(\pi' \| \pi) = \frac{1}{n} \left| \left\{ t \in [n] : \frac{\pi'(y^t | x^t)}{\pi(y^t | x^t)} \geq N \right\} \right|, \quad \text{Cov}_N^{\pi_D}(\pi' \| \pi) = \mathbb{P}_{\pi_D} \left(\frac{\pi'(y|x)}{\pi(y|x)} \geq M \right).$$

Fix $\alpha \geq 0$ and policy $\bar{\pi}$. With probability at least $1 - \delta$, for any $\pi \in \Pi$, it holds that

$$\widehat{\text{Cov}}_N(\bar{\pi} \| \pi) \geq \frac{1}{2} \text{Cov}_{e^{2\alpha} N}^{\pi_D}(\bar{\pi} \| \pi) - \log(\mathcal{N}_\infty(\Pi, \alpha)/\delta).$$

Similarly, with probability at least $1 - \delta$, for any $\pi \in \Pi$, it holds that

$$\widehat{\text{Cov}}_N(\pi \| \bar{\pi}) \leq 2 \text{Cov}_{e^{-2\alpha} N}^{\pi_D}(\pi \| \bar{\pi}) + \log(\mathcal{N}_\infty(\Pi, \alpha)/\delta).$$

Proof of Lemma J.2. We only prove the first inequality. Let $\Pi' \subseteq \Pi$ be an α -covering of Π with $|\Pi'| = \mathcal{N}_\infty(\Pi, \alpha)$. Then, by Freedman inequality (Lemma H.3) and union bound, it holds that with probability at least $1 - \delta$, for any $\pi' \in \Pi'$,

$$\widehat{\text{Cov}}_{e^\alpha N}(\bar{\pi} \| \pi') \geq \frac{1}{2} \text{Cov}_{e^\alpha N}^{\pi_D}(\bar{\pi} \| \pi') - \log(|\Pi'|/\delta).$$

Then, note that for any $\pi \in \Pi$, there exists $\pi' \in \Pi'$ such that $|\log \pi(y|x) - \log \pi'(y|x)| \leq \alpha$ for $\forall x, y$, we know

$$\left\{ t \in [n] : \frac{\bar{\pi}(y^t | x^t)}{\pi'(y^t | x^t)} \geq e^\alpha N \right\} \subseteq \left\{ t \in [n] : \frac{\bar{\pi}(y^t | x^t)}{\pi(y^t | x^t)} \geq N \right\}$$

and hence $\widehat{\text{Cov}}_{e^\alpha N}(\bar{\pi} \| \pi') \leq \widehat{\text{Cov}}_N(\bar{\pi} \| \pi)$. Similarly, $\text{Cov}_{e^\alpha N}^{\pi_D}(\bar{\pi} \| \pi') \geq \text{Cov}_{e^{2\alpha} N}^{\pi_D}(\bar{\pi} \| \pi)$. Hence, under the above event, it holds that

$$\begin{aligned} \widehat{\text{Cov}}_N(\bar{\pi} \| \pi) &\geq \widehat{\text{Cov}}_{e^\alpha N}(\bar{\pi} \| \pi') \geq \frac{1}{2} \text{Cov}_{e^\alpha N}^{\pi_D}(\bar{\pi} \| \pi') - \log(|\Pi'|/\delta) \\ &\geq \frac{1}{2} \text{Cov}_{e^{2\alpha} N}^{\pi_D}(\bar{\pi} \| \pi) - \log(|\Pi'|/\delta). \end{aligned}$$

Since $\pi \in \Pi$ is arbitrary, the proof is hence completed. $\square$

J.3 PROOF OF THEOREM F.1

By definition and concavity of $\theta \mapsto \pi(y | x)$, we know θ^* is the optimal solution of the following convex problem

$$\theta^* = \arg \min_{\theta \in \Theta} -\mathbb{E}_{(x,y) \sim \pi_0} [\log \pi_\theta(y | x)].$$

Hence, the optimality of θ^* implies

$$\langle \theta - \theta^*, -\mathbb{E}_{\pi_0} [\nabla \log \pi_{\theta^*}(y | x)] \rangle \geq 0, \quad \forall \theta \in \Theta.$$

Consider the function $F(\theta) = \mathbb{E}_{\pi_0} \left[\frac{\pi_\theta(y|x)}{\pi_{\theta^*}(y|x)} \right] - 1$, which is also convex by [Assumption F.1](#). Further, for any $\theta \in \Theta$,

$$\langle \theta - \theta^*, -\nabla F(\theta^*) \rangle = \left\langle \theta - \hat{\theta}, -\mathbb{E}_{\pi_0} \left[\frac{\nabla \pi_{\theta^*}(y | x)}{\pi_{\theta^*}(y | x)} \right] \right\rangle = \left\langle \theta - \hat{\theta}, -\mathbb{E}_{\pi_0} [\nabla \log \pi_{\theta^*}(y | x)] \right\rangle \geq 0.$$

Therefore, F attains its maximum over Θ at θ^* , i.e., $F(\theta) \leq F(\theta^*)$ for any $\theta \in \Theta$.

Similarly, under [Assumption F.1](#), it is clear that $\theta \mapsto \sum_{i=1}^n \log \pi_\theta(y^i | x^i)$ is concave, and hence $\hat{\pi} = \pi_{\hat{\theta}}$, where $\hat{\theta} \in \Theta$ satisfies

$$\left\langle \theta - \hat{\theta}, \sum_{i=1}^n -\nabla \log \pi_{\hat{\theta}}(y^i | x^i) \right\rangle \geq 0, \quad \forall \theta \in \Theta.$$

In particular, we consider the function

$$\hat{F}(\theta) := \sum_{i=1}^n \left[\frac{\pi_\theta(y^i | x^i)}{\pi_{\hat{\theta}}(y^i | x^i)} - 1 \right].$$

By definition, $\hat{F}$ is concave, and for any $\theta \in \Theta$,

$$\left\langle \theta - \hat{\theta}, -\nabla \hat{F}(\hat{\theta}) \right\rangle = \left\langle \theta - \hat{\theta}, -\sum_{i=1}^n \frac{\nabla \pi_{\hat{\theta}}(y^i | x^i)}{\pi_{\hat{\theta}}(y^i | x^i)} \right\rangle = \left\langle \theta - \hat{\theta}, \sum_{i=1}^n -\nabla \log \pi_{\hat{\theta}}(y^i | x^i) \right\rangle \geq 0.$$

Therefore, $\hat{F}$ attains its maximum over Θ at $\hat{\theta}$, and in particular, $\hat{F}(\theta^*) \leq \hat{F}(\hat{\theta}) = 0$. This implies

$$\sum_{i=1}^n \left[\frac{\pi_{\theta^*}(y^i | x^i)}{\hat{\pi}(y^i | x^i)} - \log \frac{\pi_{\theta^*}(y^i | x^i)}{\hat{\pi}(y^i | x^i)} - 1 \right] \leq \sum_{i=1}^n \log \hat{\pi}(y^i | x^i) - \sum_{i=1}^n \log \pi_{\theta^*}(y^i | x^i). \quad (52)$$

In the following, we fix any $N \geq 2$. Note that $x - \log x - 1 \geq 0$ for any $x > 0$, and $x \mapsto x - \log x - 1$ is increasing for $x \geq 1$. Therefore, (52) implies that

$$(N - \log N - 1) \cdot n \cdot \widehat{\text{Cov}}_N(\pi_{\theta^*} \| \hat{\pi}) \leq \hat{L}_n(\hat{\pi}) - \hat{L}_n(\pi_{\theta^*}). \quad (53)$$

Then, by [Lemma J.2](#), we have with probability at least $1 - \delta$, for all $\pi \in \Pi$,

$$\widehat{\text{Cov}}_N(\pi_{\theta^*} \| \pi) \geq \frac{1}{2} \cdot \mathbb{P}_{\pi_0} \left(\frac{\pi_{\theta^*}(y | x)}{\pi(y | x)} \geq e^{2\alpha} N \right) - \frac{\log(\mathcal{N}_\infty(\Pi, \alpha)/\delta)}{n}, \quad \forall \pi \in \Pi.$$

Further, by [Lemma H.4](#), the following holds with probability at least $1 - \delta$: For any $\theta \in \Theta$,

$$\begin{aligned} \hat{L}_n(\pi_\theta) - \hat{L}_n(\pi_{\theta^*}) &= \sum_{i=1}^n \log \frac{\pi_\theta(y^i | x^i)}{\pi_{\theta^*}(y^i | x^i)} \\ &\leq n \log \mathbb{E}_{\pi_0} \left[\frac{\pi_\theta(y | x)}{\pi_{\theta^*}(y | x)} \right] + \inf_{\epsilon \geq 0} \{ \log(\mathcal{N}_\infty(\Pi, \epsilon)/\delta) + 2n\epsilon \} \\ &\leq \inf_{\epsilon \geq 0} \{ \log(\mathcal{N}_\infty(\Pi, \epsilon)/\delta) + 2n\epsilon \}, \end{aligned}$$

where we use $\mathbb{E}_{\pi_0} \left[\frac{\pi_\theta(y|x)}{\pi_{\theta^*}(y|x)} \right] = F(\theta) + 1 \leq 1$ for any $\theta \in \Theta$. By union bound, we have shown that with probability at least $1 - 2\delta$,

$$\mathbb{P}_{\pi_0} \left(\frac{\pi_{\theta^*}(y | x)}{\hat{\pi}(y | x)} \geq e^{2\alpha} N \right) \lesssim \frac{\log(\mathcal{N}_\infty(\Pi, \alpha)/\delta)}{n} + \frac{1}{N} \inf_{\epsilon \geq 0} \left\{ \frac{\log \mathcal{N}_\infty(\Pi, \epsilon)}{n} + \epsilon \right\}.$$

Note that

$$\begin{aligned}\text{Cov}_{e^{2\alpha}NN'}(\hat{\pi}) &= \mathbb{P}_{\pi_0} \left(\frac{\pi_D(y|x)}{\hat{\pi}(y|x)} \geq e^{2\alpha}NN' \right) \\ &\leq \mathbb{P}_{\pi_0} \left(\frac{\pi_{\theta^*}(y|x)}{\pi(y|x)} \geq e^{2\alpha}N \right) + \mathbb{P}_{\pi_0} \left(\frac{\pi_D(y|x)}{\pi_{\theta^*}(y|x)} \geq N' \right).\end{aligned}$$

Therefore, the proof is completed by rescaling $N \leftarrow Ne^{-2\alpha}/N'$ and combining the inequalities above.

J.4 PROOFS FOR SUPPORTING RESULTS

Proof of Proposition F.1 (a). Assume that $B \geq \frac{1}{2} \log(4n)$ and $n \geq d$. Consider $\mathcal{X} = \perp$, $\mathcal{Y} = [d]$ and the feature map be given by $\phi(y) = Be_y$ for $y \in \mathcal{Y}$, where $(e_1, \dots, e_d)$ is the coordinate basis of $\mathbb{R}^d$. For the simplicity of our argument, we consider $\Theta = \{\theta \in \mathbb{R}^d : \|\theta\|_\infty \leq 1\}$, and we set

$$\theta^* = \frac{\log(4n)}{2B} \cdot \left(e_1 - \sum_{j=2}^d e_j \right)$$

Then it holds that

$$\pi_D(1) = \frac{4n}{d-1+4n}, \quad \pi_D(y) = \frac{1}{d-1+4n}, \quad \forall y > 1.$$

Therefore, under $\mathcal{D} \sim \pi_D$, it holds that

$$\mathbb{E} \left[\sum_{t=1}^n \mathbb{I}\{y^t \neq 1\} \right] \leq \frac{n(d-1)}{d-1+4n} \leq \frac{d-1}{4} \leq \frac{n}{2}.$$

In particular, with probability at least 0.5, it holds that $\sum_{t=1}^T \mathbb{I}\{y^t \neq 1\} \leq \frac{d-1}{2}$. i.e., the set $\mathcal{Y}_D = \mathcal{Y} \setminus \mathcal{D}$ has cardinality at least $\frac{d}{2}$.

In the following, we condition on this event and analyze the MLE $\hat{\theta}$. By the definition of MLE, for any $y \in \mathcal{Y}_D$, it must hold that $\hat{\theta}_y = -1$, and we also know $\hat{\theta}_1 = 1$. This implies $\pi_{\hat{\theta}}(y) \leq \frac{1}{e^{2B}}$ for any $y \in \mathcal{Y}_D$. Therefore, for $N \leq \frac{e^B}{4n+d-1}$, we have

$$\text{Cov}_N(\pi_{\hat{\theta}}) \geq \pi_D(\mathcal{Y}_D) \geq \frac{d}{2(d-1+4n)} \geq \frac{d}{10n}.$$

This is the desired lower bound. $\square$

Proof of Proposition F.1 (b). Let $\epsilon = c_0 \sqrt{\frac{d}{n}}$ and $p = \frac{\epsilon^2}{\log N}$ for a sufficiently small absolute constant $c_0 > 0$. Let $\mathcal{X} = \{0, 1, \dots, d\}$, $\mathcal{Y} = \{0, 1\}$, and the distribution μ be given by $\mu(0) = p$, $\mu(1) = \dots = \mu(d) = \frac{1-p}{d}$.

Consider $\pi_D(\cdot | i) = \text{Ber}(1/2)$ for $i \in [d]$ and $\pi_D(1|0) = 1$. For any $\theta \in \Theta := \{+1, -1\}^d$, we define π_θ as

$$\pi_\theta(1|0) = \frac{1}{N}, \quad \pi_\theta(\cdot|i) = \text{Ber}\left(\frac{1+\epsilon\theta_i}{2}\right), \quad \forall i \in [d].$$

Then, we can calculate

$$\begin{aligned}& \max_{\theta \in \Theta} \hat{L}_n(\pi_\theta) - \hat{L}_n(\pi_D) \\ &= -N(0) \log N + \frac{1}{2} \sum_{i \in [d]} \left[|N(i, +) - N(i, -)| \log \frac{1+\epsilon}{1-\epsilon} + N(i) \log(1-\epsilon^2) \right] \\ &\geq -N(0) \log N - n\epsilon^2 + \frac{\epsilon}{2} \sum_{i \in [d]} |N(i, 0) - N(i, 1)|,\end{aligned}$$

where we denote $N(x, y) = \#\{t \in [n] : (x^t, y^t) = (x, y)\}$, $N(x) = N(x, 0) + N(x, 1)$. In the following, we denote $\Delta_i = N(i, 1) - N(i, 0)$. Note that Δ_i is a sum of n i.i.d $\{-1, 0, 1\}$ -valued random variables with mean zero and variance $\frac{1-p}{d}$, and hence $\mathbb{P}_{\pi_0}(|\Delta_i| \geq c\sqrt{\frac{n}{d}}) \geq 0.99$ for an absolute constant $c > 0$. Then, it is clear that

$$\mathbb{P}_{\pi_0}\left(\sum_{i=1}^d |\Delta_i| \leq \frac{d}{2} \cdot c\sqrt{\frac{n}{d}}\right) \leq \frac{1}{4}.$$

By Markov inequality, we also know $\mathbb{P}_{\pi_0}(N(0) \geq 4np) \leq \frac{1}{4}$. Combining the inequalities above, we know with probability at least 0.5, we have

$$\max_{\theta \in \Theta} L(\pi_\theta) - L(\pi_0) \geq -4np \log N - n\epsilon^2 + \frac{c\epsilon\sqrt{nd}}{4} > 0.$$

This implies that the MLE $\hat{\pi} \in \{\pi_\theta\}_{\theta \in \Theta}$, and hence $\text{Cov}_N(\hat{\pi}) \geq p$. This is the desired lower bound. $\square$

K PROOFS FROM SECTION 5

Organization. We begin with the proof of [Proposition 5.1](#) (the upper bound), which is relatively simple and serves as motivation. We then present the proofs of [Theorem 4.2](#) and [Theorem 5.1](#), which are more involved. Finally, the proofs of the lower bounds are given in the remaining subsections.

Notation. For notational simplicity, we denote

$$\bar{\phi}_\theta(x, y_{1:h-1}) = \mathbb{E}_{y_h \sim \pi_\theta(\cdot | x, y_{1:h-1})}[\phi(x, y_{1:h})],$$

and

$$\begin{aligned} \phi^*(x, y_{1:h}) &:= \phi(x, y_{1:h}) - \bar{\phi}_{\theta^*}(x, y_{1:h-1}), \\ \text{Var}_{\pi_0}(x, y_{1:h-1}) &:= \mathbb{E}_{y_h \sim \pi_0(\cdot | x, y_{1:h-1})} \|\phi^*(x, y_{1:h})\|^2. \end{aligned}$$

Then, by definition,

$$\nabla \log \pi_\theta(y_{1:H} | x) = \sum_{h=1}^H \phi^*(x, y_{1:h}) + \sum_{h=1}^H (\bar{\phi}_{\theta^*}(x, y_{1:h-1}) - \bar{\phi}_\theta(x, y_{1:h-1})). \quad (54)$$

We also write

$$\epsilon_\theta(x, y_{1:h-1}) = D_{\text{KL}}(\pi_0(\cdot | x, y_{1:h-1}) \| \pi_\theta(\cdot | x, y_{1:h-1})).$$

By concavity, we have

$$\epsilon_\theta(x, y_{1:h-1}) \leq \langle \bar{\phi}_\theta(x, y_{1:h-1}) - \bar{\phi}_{\theta^*}(x, y_{1:h-1}), \theta - \theta^* \rangle. \quad (55)$$

By [Lemma H.5](#), it holds that

$$\|\bar{\phi}_{\theta^*}(x, y_{1:h-1}) - \bar{\phi}_\theta(x, y_{1:h-1})\| \leq 2\sqrt{\text{Var}_{\pi_0}(x, y_{1:h-1}) \cdot \epsilon_\theta(x, y_{1:h-1})} + 3B\epsilon_\theta(x, y_{1:h-1}). \quad (56)$$

For notational simplicity, for any $f : \mathcal{X} \times \mathcal{A}^* \rightarrow \mathbb{R}$ and dataset $\mathcal{D} = \{(x^i, y_{1:H}^i)\}_{i \in [n]}$, we write

$$\widehat{\mathbb{E}}_{\mathcal{D}}[f] := \frac{1}{n} \sum_{i=1}^n f(x^i, y_{1:H}^i),$$

K.1 PROOF OF PROPOSITION 5.1 (UPPER BOUND)

Because the projection operator Proj_Θ is an contraction, we have

$$\begin{aligned} &\|\theta^t - \theta^*\|^2 - \|\theta^{t+1} - \theta^*\|^2 \\ &\geq \|\theta^t - \theta^*\|^2 - \|\theta^t + \eta \nabla \log \pi_{\theta^t}(y^t | x^t) - \theta^*\|^2 \\ &= 2\eta \langle -\nabla \log \pi_{\theta^t}(y^t | x^t), \theta^t - \theta^* \rangle - 2\eta^2 \|\nabla \log \pi_{\theta^t}(y^t | x^t)\|^2. \end{aligned}$$

Telescoping and taking expectation, we have

$$\mathbb{E} \left[\sum_{t=1}^T \langle -\nabla \log \pi_{\theta^t}(y | x), \theta^t - \theta^* \rangle \right] \leq \frac{1}{2\eta} + \eta \mathbb{E} \left[\sum_{t=1}^T \mathbb{E}_{(x,y) \sim \pi_0} \|\nabla \log \pi_{\theta^t}(y | x)\|^2 \right]. \quad (57)$$

Note that $(x^t, y^t) | \theta^t \sim \pi_0$, and hence

$$\mathbb{E}[\nabla \log \pi_{\theta^t}(y^t | x^t) | \theta^t] = \mathbb{E}_{(x,y) \sim \pi_0}[\nabla \log \pi_{\theta^t}(y | x)] = \nabla_{\theta} D_{\text{KL}}(\pi_0 \| \pi_{\theta})|_{\theta=\theta^t}.$$

Further, by convexity, it holds that for any $\theta \in \Theta$,

$$G(\theta) := \mathbb{E}_{\pi_0}[\langle \nabla \log \pi_{\theta}(y | x), \theta - \theta^* \rangle] = \langle \nabla_{\theta} D_{\text{KL}}(\pi_0 \| \pi_{\theta}), \theta - \theta^* \rangle \geq D_{\text{KL}}(\pi_0 \| \pi_{\theta}).$$

Therefore, we have

$$\mathbb{E} \left[\sum_{t=1}^T D_{\text{KL}}(\pi_0 \| \pi_{\theta^t}) \right] \leq \mathbb{E} \left[\sum_{t=1}^T G(\theta^t) \right] \leq \frac{1}{2\eta} + \eta \mathbb{E} \left[\sum_{t=1}^T \mathbb{E}_{(x,y) \sim \pi_0} \|\nabla \log \pi_{\theta^t}(y | x)\|^2 \right].$$

On the other hand, using the fact that $\log \pi_{\theta}(y | x)$ is concave and (HB^2) -smooth (i.e., $-HB^2 I \preceq \nabla^2 \log \pi_{\theta}(y | x) \preceq 0$),

$$\|\nabla \log \pi_{\theta}(y | x) - \nabla \log \pi_{\theta^*}(y | x)\|^2 \leq HB^2 \cdot \langle \theta - \theta^*, \nabla \log \pi_{\theta^*}(y | x) - \nabla \log \pi_{\theta}(y | x) \rangle$$

Taking expectation of $(x, y) \sim \pi_0$ and using the fact that $\mathbb{E}_{\pi_0}[\nabla \log \pi_{\theta^*}(y | x)] = 0$, we have

$$\mathbb{E}_{\pi_0} \|\nabla \log \pi_{\theta}(y | x) - \nabla \log \pi_{\theta^*}(y | x)\|^2 \leq HB^2 \cdot G(\theta), \quad \forall \theta \in \Theta.$$

Further, note that $\mathbb{E}_{\pi_0} \|\nabla \log \pi_{\theta^*}(y | x)\|^2 = \sigma_{\star}^2$, it holds that

$$\mathbb{E}_{\pi_0} \|\nabla \log \pi_{\theta^*}(y | x)\|^2 \leq 2\sigma_{\star}^2 + 2HB^2 \cdot G(\theta), \quad \forall \theta \in \Theta. \quad (58)$$

Combining the inequalities above, we can conclude that

$$\mathbb{E} \left[\sum_{t=1}^T G(\theta^t) \right] \leq \frac{1}{2\eta} + 2\eta HB^2 \mathbb{E} \left[\sum_{t=1}^T G(\theta^t) \right] + 2\eta T \sigma_{\star}^2.$$

Therefore, as long as $\eta \leq \frac{1}{4HB^2}$, it holds

$$\frac{1}{\eta} + 4\eta T \sigma_{\star}^2 \geq \mathbb{E} \left[\sum_{t=1}^T G(\theta^t) \right] \geq \mathbb{E} \left[\sum_{t=1}^T D_{\text{KL}}(\pi_0 \| \pi_{\theta^t}) \right].$$

This is the desired upper bound. $\square$

K.2 PROOF OF THEOREM 5.1

We denote $M := \log N$, and we analyze the normalized SGD iterates assuming $\lambda \geq 3BM$ and $\frac{\lambda\eta}{M} \leq \frac{1}{8}$. and

Denote

$$\tilde{g}_{\theta}(\mathcal{D}) := \frac{\hat{g}_{\theta}(\mathcal{D})}{\lambda + \|\hat{g}_{\theta}(\mathcal{D})\|}.$$

Then the normalized SGD update can be rewritten as $\theta^{t+1} = \text{Proj}_{\Theta}(\theta + \eta \tilde{g}_{\theta^t}(\mathcal{D}))$. By the standard SGD proof, we know that

$$\sum_{t=1}^T \langle -\tilde{g}_{\theta^t}(\mathcal{D}), \theta^t - \theta^* \rangle \leq \frac{\|\theta^0 - \theta^*\|^2}{2\eta} + \eta \sum_{t=1}^T \|\tilde{g}_{\theta^t}(\mathcal{D})\|^2.$$

Taking expectation on both sides, we have

$$\mathbb{E} \left[\sum_{t=1}^T \mathbb{E}_{\mathcal{D} \sim \pi_0} \langle -\tilde{g}(\theta^t; \mathcal{D}), \theta^t - \theta^* \rangle \right] \leq \frac{\|\theta^0 - \theta^*\|^2}{2\eta} + \eta \mathbb{E} \left[\sum_{t=1}^T \mathbb{E}_{\mathcal{D} \sim \pi_0} \|\tilde{g}(\theta^t; \mathcal{D})\|^2 \right].$$

Note that $\|\tilde{g}_\theta(\mathcal{D})\| \leq \min\left\{1, \frac{\|\hat{g}_\theta(\mathcal{D})\|}{\lambda}\right\}$. In the following, we analyze $\|\hat{g}_\theta(\mathcal{D})\|$ under $\mathcal{D} = \{(x^i, y_{1:H}^i)\}_{i \in [K]} \sim \pi_0$. Recall that for notational simplicity, for any $f : \mathcal{X} \times \mathcal{A}^* \rightarrow \mathbb{R}$, we write

$$\hat{\mathbb{E}}_{\mathcal{D}}[f] := \frac{1}{K} \sum_{i=1}^K f(x^i, y_{1:H}^i),$$

which is random variable. We denote

$$\bar{g}_\theta(\mathcal{D}) := \hat{\mathbb{E}}_{\mathcal{D}} \left[\sum_{h=1}^H (\bar{\phi}_\theta(x, y_{1:h-1}) - \bar{\phi}_{\theta^*}(x, y_{1:h-1})) \right],$$

and

$$z(\mathcal{D}) := \hat{\mathbb{E}}_{\mathcal{D}} \left[\sum_{h=1}^H \phi^*(x, y_{1:h}) \right] = \hat{\mathbb{E}}_{\mathcal{D}} \left[\sum_{h=1}^H (\phi(x, y_{1:h}) - \bar{\phi}_{\theta^*}(x, y_{1:h-1})) \right].$$

Then, by definition, $-\hat{g}_\theta(\mathcal{D}) = \bar{g}_\theta(\mathcal{D}) - z(\mathcal{D})$.

Bounds on $\bar{g}_\theta(\mathcal{D})$. By (55), we know

$$\begin{aligned} \langle \bar{g}_\theta(\mathcal{D}), \theta - \theta^* \rangle &= \hat{\mathbb{E}}_{\mathcal{D}} \left[\sum_{h=1}^H \langle \bar{\phi}_\theta(x, y_{1:h-1}) - \bar{\phi}_{\theta^*}(x, y_{1:h-1}), \theta - \theta^* \rangle \right] \\ &\geq \hat{\mathbb{E}}_{\mathcal{D}} \left[\sum_{h=1}^H \epsilon_\theta(x, y_{1:h-1}) \right] =: \epsilon_\theta(\mathcal{D}). \end{aligned}$$

By (56), we also have

$$\begin{aligned} \|\bar{g}_\theta(\mathcal{D})\| &\leq \hat{\mathbb{E}}_{\mathcal{D}} \left[\sum_{h=1}^H \|\bar{\phi}_\theta(x, y_{1:h-1}) - \bar{\phi}_{\theta^*}(x, y_{1:h-1})\| \right] \\ &\leq \hat{\mathbb{E}}_{\mathcal{D}} \left[\sum_{h=1}^H 2\sqrt{\text{Var}_{\pi_0}(x, y_{1:h-1}) \cdot \epsilon_\theta(x, y_{1:h-1})} + 3B\epsilon_\theta(x, y_{1:h-1}) \right] \\ &\leq 2\sqrt{\sigma^2(\mathcal{D}) \cdot \epsilon_\theta(\mathcal{D})} + 3B\epsilon_\theta(\mathcal{D}), \end{aligned}$$

where we denote

$$\sigma^2(\mathcal{D}) := \hat{\mathbb{E}}_{\mathcal{D}} \left[\sum_{h=1}^H \text{Var}_{\pi_0}(x, y_{1:h-1}) \right].$$

Bounds on $z(\mathcal{D})$. Note that $K \cdot z(\mathcal{D}) = K \cdot \hat{\mathbb{E}}_{\mathcal{D}} \left[\sum_{h=1}^H \phi^*(x, y_{1:h}) \right] = \sum_{i=1}^K \sum_{h=1}^H \phi^*(x^i, y_{1:h}^i)$ is a sum of the martingale difference sequence $\{\phi^*(x^i, y_{1:h}^i)\}_{i \in [K], h \in [H]}$. Therefore, we can calculate

$$\mathbb{E}_{\pi_0} \|z(\mathcal{D})\|^2 = \frac{1}{K} \mathbb{E}_{\pi_0} \left[\sum_{h=1}^H \|\phi^*(x, y_{1:h})\|^2 \right] = \frac{\sigma_*^2}{K}.$$

Furthermore, by Freedman's inequality (Lemma H.2), for any fixed vector v , parameter $\gamma \in (0, \frac{1}{B})$ and $\delta \in (0, 1)$, it holds that

$$\mathbb{P}_{\pi_0} \left(\sum_{i=1}^K \sum_{h=1}^H (\langle \phi^*(x^i, y_{1:h}^i), v \rangle - \gamma \mathbb{E}[\langle \phi^*(x^i, y_{1:h}^i), v \rangle^2 \mid x^i, y_{1:h-1}^i]) \geq \gamma^{-1} \log(1/\delta) \right) \leq \delta.$$

Note that for $v = \theta - \theta^*$, by Lemma H.6, we have

$$\begin{aligned} &\mathbb{E}[\langle \phi^*(x^i, y_{1:h}^i), v \rangle^2 \mid x^i, y_{1:h-1}^i] \\ &= \mathbb{E}_{y_h \sim \pi_0(\cdot \mid x^i, y_{1:h-1}^i)} \langle \phi^*(x^i, y_{1:h-1}^i, y_h), v \rangle^2 \\ &\leq 15BD_{\text{KL}}(\pi_0(\cdot \mid x^i, y_{1:h-1}^i) \parallel \pi_\theta(\cdot \mid x^i, y_{1:h-1}^i)) = 15B\epsilon_\theta(x^i, y_{1:h-1}^i). \end{aligned}$$

Therefore, setting $\gamma = \frac{1}{30B}$, we have shown that for any $\delta \in (0, 1)$, it holds that

$$\mathbb{P}_{\pi_0} \left(\langle z(\mathcal{D}), \theta - \theta^* \rangle \geq \frac{1}{2} \widehat{\mathbb{E}}_{\mathcal{D}} \left[\sum_{h=1}^H \epsilon_{\theta}(x, y_{1:h-1}) \right] + \frac{30B \log(1/\delta)}{K} \right) \leq \delta.$$

Recall that we denote $\epsilon_{\theta}(\mathcal{D}) := \widehat{\mathbb{E}}_{\mathcal{D}} \left[\sum_{h=1}^H \epsilon_{\theta}(x, y_{1:h-1}) \right]$. Therefore, taking integration gives

$$\mathbb{E}_{\pi_0} \left(\langle z(\mathcal{D}), \theta - \theta^* \rangle - \frac{1}{2} \epsilon_{\theta}(\mathcal{D}) \right)_+ \leq \frac{30B}{K} =: \alpha_K.$$

Upper bounding $\|\tilde{g}_{\theta}(\mathcal{D})\|$. Using $\|\tilde{g}_{\theta}(\mathcal{D})\| \leq \min\left\{1, \frac{\|\widehat{g}_{\theta}(\mathcal{D})\|}{\lambda}\right\}$, we know

$$\|\tilde{g}_{\theta}(\mathcal{D})\|^2 \leq \mathbb{I}\{\epsilon_{\theta}(\mathcal{D}) \geq M\} + \mathbb{I}\{\epsilon_{\theta}(\mathcal{D}) \leq M\} \cdot \frac{\|\widehat{g}_{\theta}(\mathcal{D})\|}{\lambda}.$$

Therefore, for any fixed θ , it holds that

$$\begin{aligned} & \mathbb{E}_{\pi_0} \|\tilde{g}_{\theta}(\mathcal{D})\|^2 \\ & \leq \mathbb{P}_{\pi_0}(\epsilon_{\theta}(\mathcal{D}) \geq M) + \frac{1}{\lambda} \mathbb{E}_{\pi_0} [\mathbb{I}\{\epsilon_{\theta}(\mathcal{D}) \leq M\} \cdot \|\widehat{g}_{\theta}(\mathcal{D})\|] + \frac{1}{\lambda} \mathbb{E}_{\pi_0} \|z(\mathcal{D})\| \\ & \leq \max\left\{\frac{1}{M}, \frac{3B}{\lambda}\right\} \cdot \mathbb{E}_{\pi_0} \min\{M, \epsilon_{\theta}(\mathcal{D})\} + \frac{2}{\lambda} \mathbb{E}_{\pi_0} \sqrt{\sigma^2(\mathcal{D}) \cdot \min\{M, \epsilon_{\theta}(\mathcal{D})\}} + \frac{\sigma_{\star}}{\lambda\sqrt{K}} \\ & \leq \frac{1}{M} \mathbb{E}_{\pi_0} \min\{M, \epsilon_{\theta}(\mathcal{D})\} + \frac{\sigma_{\star}}{\lambda} \left[2\sqrt{\mathbb{E}_{\pi_0} \min\{M, \epsilon_{\theta}(\mathcal{D})\}} + \frac{1}{\sqrt{K}} \right], \end{aligned}$$

where the last inequality follows from $\lambda \geq 3BM$, $\mathbb{E}[\sigma^2(\mathcal{D})] = \sigma_{\star}^2$, and Cauchy's inequality.

Lower bounding $\langle -\tilde{g}_{\theta}(\mathcal{D}), \theta - \theta^* \rangle$. By the inequalities above, we know

$$\begin{aligned} \Lambda_{\theta} &:= \mathbb{E}_{\pi_0} \langle -\tilde{g}_{\theta}(\mathcal{D}), \theta - \theta^* \rangle \\ &= \mathbb{E}_{\pi_0} \left[\frac{\langle \tilde{g}_{\theta}(\mathcal{D}), \theta - \theta^* \rangle - \langle z(\mathcal{D}), \theta - \theta^* \rangle}{\lambda + \|\widehat{g}_{\theta}(\mathcal{D})\|} \right] \\ &\geq \mathbb{E}_{\pi_0} \left[\frac{\epsilon_{\theta}(\mathcal{D}) - \langle z(\mathcal{D}), \theta - \theta^* \rangle}{\lambda + \|\widehat{g}_{\theta}(\mathcal{D})\|} \right] \\ &\geq \frac{1}{2} \mathbb{E}_{\pi_0} \left[\frac{\epsilon_{\theta}(\mathcal{D})}{\lambda + \|\widehat{g}_{\theta}(\mathcal{D})\|} \right] - \frac{1}{\lambda} \mathbb{E}_{\pi_0} \left[\left(\langle z(\mathcal{D}), \theta - \theta^* \rangle - \frac{1}{2} \epsilon_{\theta}(\mathcal{D}) \right)_+ \right] \\ &\geq \frac{1}{2} \mathbb{E}_{\pi_0} \left[\frac{\epsilon_{\theta}(\mathcal{D})}{\lambda + \|z(\mathcal{D})\| + \|\widehat{g}_{\theta}(\mathcal{D})\|} \right] - \frac{\alpha_K}{\lambda}. \end{aligned}$$

Note that

$$\begin{aligned} & \lambda + \|z(\mathcal{D})\| + \|\widehat{g}_{\theta}(\mathcal{D})\| \\ & \leq \lambda + \|z(\mathcal{D})\| + 2\sqrt{\sigma^2(\mathcal{D}) \cdot \epsilon_{\theta}(\mathcal{D})} + 3B\epsilon_{\theta}(\mathcal{D}) \\ & \leq \frac{\max\{M, \epsilon_{\theta}(\mathcal{D})\}}{M} \cdot \left[2\lambda + \|z(\mathcal{D})\| + 2M\sqrt{\frac{\sigma^2(\mathcal{D})}{\min\{M, \epsilon_{\theta}(\mathcal{D})\}}} \right], \end{aligned}$$

where we use $\min\{M, x\} \max\{M, x\} = Mx$, and $\lambda \geq 3BM$. Combining these two inequalities, we have

$$\begin{aligned} 2\Lambda_{\theta} + \frac{2\alpha_K}{\lambda} &\geq \mathbb{E}_{\pi_0} \left[\frac{\epsilon_{\theta}(\mathcal{D})}{\lambda + \|z(\mathcal{D})\| + \|\widehat{g}_{\theta}(\mathcal{D})\|} \right] \\ &\geq \mathbb{E}_{\pi_0} \left[\frac{\min\{M, \epsilon_{\theta}(\mathcal{D})\}}{2\lambda + \|z(\mathcal{D})\| + 2M\sqrt{\sigma^2(\mathcal{D})/\min\{M, \epsilon_{\theta}(\mathcal{D})\}}} \right] \\ &\geq \frac{(\mathbb{E}_{\pi_0} \min\{M, \epsilon_{\theta}(\mathcal{D})\})^2}{\mathbb{E}_{\pi_0} [\min\{M, \epsilon_{\theta}(\mathcal{D})\} (2\lambda + \|z(\mathcal{D})\|)] + 2M\sqrt{\mathbb{E}_{\pi_0} \sigma^2(\mathcal{D}) \cdot \mathbb{E}_{\pi_0} \min\{M, \epsilon_{\theta}(\mathcal{D})\}}} \\ &\geq \frac{(\mathbb{E}_{\pi_0} \min\{M, \epsilon_{\theta}(\mathcal{D})\})^2}{2\lambda \mathbb{E}_{\pi_0} \min\{M, \epsilon_{\theta}(\mathcal{D})\} + M\sqrt{\sigma_{\star}^2/K} + 2M\sqrt{\sigma_{\star}^2 \mathbb{E}_{\pi_0} \min\{M, \epsilon_{\theta}(\mathcal{D})\}}}, \end{aligned}$$

where the last two inequalities follow from Cauchy's inequality.

Putting everything together. Denote $\Delta_\theta := \mathbb{E}_{\pi_0} \min\{M, \epsilon_\theta(\mathcal{D})\}$. Note that by [Proposition D.10](#), we have (recall $M := \log N$)

$$\begin{aligned} \text{Cov}_N(\pi_\theta) &\leq \frac{1}{M} \mathbb{E}_{\pi_0} \min\left\{M, \sum_{h=1}^H \epsilon_\theta(x, y_{1:h-1})\right\} \\ &\leq \frac{1}{M} \mathbb{E}_{\mathcal{D} \sim \pi_0} \min\left\{M, \widehat{\mathbb{E}}_{\mathcal{D}} \left[\sum_{h=1}^H \epsilon_\theta(x, y_{1:h-1}) \right]\right\} = \frac{1}{M} \Delta_\theta. \end{aligned}$$

Then, we have shown that for any fixed parameter θ ,

$$\mathbb{E}_{\pi_0} \|\tilde{g}_\theta(\mathcal{D})\|^2 \leq \frac{1}{M} \Delta_\theta + \frac{\sigma_\star}{\lambda} \left[2\sqrt{\Delta_\theta} + \frac{1}{\sqrt{K}} \right],$$

and

$$\Lambda_\theta = \mathbb{E}_{\pi_0} \langle -\tilde{g}_\theta(\mathcal{D}), \theta - \theta^\star \rangle \geq \frac{1}{2} \frac{\Delta_\theta^2}{2\lambda\Delta_\theta + M\sigma_\star \left[\frac{1}{\sqrt{K}} + \sqrt{\Delta_\theta} \right]} - \frac{\alpha_K}{\lambda}.$$

Finally, note that implies that

$$\begin{aligned} \mathbb{E} \left[\sum_{t=1}^T \Lambda_{\theta^t} \right] &\leq \frac{1}{2\eta} + \eta \mathbb{E} \left[\sum_{t=1}^T \mathbb{E}_{\mathcal{D} \sim \pi_0} \|\tilde{g}(\theta^t; \mathcal{D})\|^2 \right] \\ &\leq \frac{1}{2\eta} + \frac{\eta}{M} \mathbb{E} \left[\sum_{t=1}^T \Delta_{\theta^t} \right] + \frac{\sigma_\star}{\lambda} \left[2\mathbb{E} \left[\sum_{t=1}^T \sqrt{\Delta_{\theta^t}} \right] + \frac{T}{\sqrt{K}} \right]. \end{aligned}$$

Therefore, we define $\Delta = \frac{1}{T} \mathbb{E} \left[\sum_{t=1}^T \Delta_{\theta^t} \right]$, and then by Cauchy inequality,

$$\begin{aligned} \frac{1}{2T\eta} + \frac{\eta}{M} \Delta + \frac{\eta\sigma_\star}{\lambda} \left[2\sqrt{\Delta} + \frac{1}{\sqrt{K}} \right] &\geq \frac{1}{T} \mathbb{E} \left[\sum_{t=1}^T \Lambda_{\theta^t} \right] \\ &\geq \frac{1}{2T} \mathbb{E} \left[\sum_{t=1}^T \frac{\Delta_{\theta^t}^2}{2\lambda\Delta_{\theta^t} + M\sigma_\star \left[\frac{1}{\sqrt{K}} + \sqrt{\Delta_{\theta^t}} \right]} \right] - \frac{\alpha_K}{\lambda} \\ &\geq \frac{1}{2} \frac{\Delta^2}{2\lambda\Delta + M\sigma_\star \left[\frac{1}{\sqrt{K}} + \sqrt{\Delta} \right]} - \frac{\alpha_K}{\lambda}. \end{aligned}$$

Re-organizing, we know as long as $\frac{\lambda\eta}{M} \leq \frac{1}{8}$, it holds that

$$\Delta \lesssim \left(\frac{M\sigma_\star}{T\eta} \right)^{2/3} + \frac{\lambda}{T\eta} + (\eta\sigma_\star)^2 + \frac{\eta M\sigma_\star^2}{\lambda} + \left(\frac{M\sigma_\star}{\lambda K} \right)^{2/3} + \frac{B}{K}.$$

In particular, we choose $\lambda = \frac{M}{8\eta}$ and require $\eta \leq \frac{1}{24B}$, we have

$$\Delta \lesssim \left(\frac{M\sigma_\star}{T\eta} \right)^{2/3} + \frac{M}{T\eta^2} + (\eta\sigma_\star)^2 + \frac{B}{K}.$$

Choosing $\eta = \min \left\{ \frac{1}{24B}, \left(\frac{M}{\sigma_\star^2 T} \right)^{1/4} \right\}$, we have

$$\Delta \lesssim \sqrt{\frac{\sigma_\star^2 M}{T}} + \frac{B^2 M}{T} + \frac{B}{K},$$

which implies

$$\frac{1}{T} \mathbb{E} \left[\sum_{t=1}^T \text{Cov}_N(\pi_{\theta^t}) \right] \leq \frac{1}{T} \mathbb{E} \left[\sum_{t=1}^T \frac{1}{M} \Delta_{\theta^t} \right] \leq \sqrt{\frac{\sigma_\star^2}{TM}} + \frac{B^2}{T} + \frac{B}{KM}.$$

This is the desired upper bound. $\square$

Remark K.1 (Comparison to standard convex optimization analyses). *On a technical level, we find the proof of [Theorem 5.1](#) to be interesting because it does not pass through KL divergence as an intermediate quantity. More broadly, we do not know how to derive the result as an application of standard analysis techniques in optimization (e.g., via a gradient dominance or PL-type condition), but it would be interesting to see if there is a connection.*⁸

K.3 PROOF OF [THEOREM 4.2](#)

We prove the following slightly stronger result. [Theorem 4.2](#) follows immediately by combining [Theorem K.1](#) and [Proposition D.10](#).

Theorem K.1. *Suppose that [Assumption 2.2](#) holds. Then the MLE $\hat{\pi}$ achieves*

$$\mathbb{E}_{\mathcal{D}}[D_{\text{seq},N}(\pi_{\mathcal{D}} \parallel \hat{\pi})] \lesssim \sqrt{\frac{\sigma_{\star}^2 \log N}{n}} + \frac{B^2 \log N}{n},$$

for any parameter $N \geq 1$, where the divergence $D_{\text{seq},N}(\cdot \parallel \cdot)$ is defined in [Proposition D.10](#).

The following lemma follows from the optimality of the MLE $\hat{\pi} = \pi_{\hat{\theta}}$, i.e.,

$$\hat{\theta} = \arg \max_{\theta \in \Theta} \mathbb{E}_{\mathcal{D}}[\log \pi_{\theta}(y_{1:H} \mid x)].$$

Lemma K.1. *Denote*

$$E_1 := \mathbb{E}_{\mathcal{D}} \left[\sum_{h=1}^H \epsilon_{\hat{\theta}}(x, y_{1:h-1}) \right] = \mathbb{E}_{\mathcal{D}} \left[\sum_{h=1}^H D_{\text{KL}}(\pi_{\mathcal{D}}(\cdot \mid x, y_{1:h-1}) \parallel \hat{\pi}(\cdot \mid x, y_{1:h-1})) \right]. \quad (59)$$

Then it holds that $\mathbb{E}[E_1] \leq \frac{2\sigma_{\star}}{\sqrt{n}}$.

In the following, we prove concentration bounds on E_1 . For simplicity, we denote $A = \log N$.

Lemma K.2. *Fix any $\Delta \in (0, \frac{1}{100B}]$, $\delta \in (0, 1)$, and let $J = \exp(\frac{1}{\Delta^2} + 2) \log(1/\delta)$. Let $\Theta' := \{\theta_1, \dots, \theta_J\}$, where $\theta_1, \dots, \theta_J \sim \mathcal{N}(0, \Delta^2 I)$ are sampled i.i.d. Then the following holds with probability at least $1 - \delta$ over the randomness of Θ' and $\mathcal{D}$:*

(1) *For any $j \in [J]$, it holds that*

$$\mathbb{E}_{\pi_0} \min \left\{ A, \sum_{h=1}^H \epsilon_{\theta_j}(x, y_{1:h-1}) \right\} \leq 2\mathbb{E}_{\mathcal{D}} \min \left\{ A, \sum_{h=1}^H \epsilon_{\theta_j}(x, y_{1:h-1}) \right\} + \frac{8A \log(4J/\delta)}{n}.$$

(2) *There exists $j \in [J]$ such that*

$$\mathbb{E}_{\pi_0} \min \left\{ A, \sum_{h=1}^H \epsilon_{\hat{\theta}}(x, y_{1:h-1}) \right\} \leq 2\mathbb{E}_{\pi_0} \min \left\{ A, \sum_{h=1}^H \epsilon_{\theta_j}(x, y_{1:h-1}) \right\} + 100\Delta^2 \sigma_{\star}^2, \quad (60)$$

and

$$\begin{aligned} \mathbb{E}_{\mathcal{D}} \min \left\{ A, \sum_{h=1}^H \epsilon_{\theta_j}(x, y_{1:h-1}) \right\} &\leq 2\mathbb{E}_{\mathcal{D}} \min \left\{ A, \sum_{h=1}^H \epsilon_{\hat{\theta}}(x, y_{1:h-1}) \right\} \\ &\quad + 100\Delta^2 \mathbb{E}_{\mathcal{D}} \left[\sum_{h=1}^H \text{Var}_{\pi_0}(x, y_{1:h-1}) \right]. \end{aligned} \quad (61)$$

⁸We further note that the inherent variance $\sigma_{\star}^2$ corresponds to the gradient variance at the true parameter θ^* , and hence is tighter than typical analyses that depend on global notions of variance.

Now, we condition on the success event $\mathcal{E}$ of [Lemma K.2](#), and let $j \in [J]$ be an index such that (60) and (61) hold. Then, we can upper bound (recall that $A = \log N$)

$$\begin{aligned}
D_{\text{seq},N}(\pi_D \parallel \pi_{\hat{\theta}}) &= \mathbb{E}_{\pi_D} \min \left\{ A, \sum_{h=1}^H \epsilon_{\hat{\theta}}(x, y_{1:h-1}) \right\} \\
&\leq 2 \mathbb{E}_{\pi_D} \min \left\{ A, \sum_{h=1}^H \epsilon_{\theta_j}(x, y_{1:h-1}) \right\} + 100\Delta^2\sigma_{\star}^2 \\
&\leq 4\hat{\mathbb{E}}_{\mathcal{D}} \min \left\{ A, \sum_{h=1}^H \epsilon_{\theta_j}(x, y_{1:h-1}) \right\} + \frac{16A \log(4J/\delta)}{n} + 100\Delta^2\sigma_{\star}^2 \\
&\leq 8\hat{\mathbb{E}}_{\mathcal{D}} \min \left\{ A, \sum_{h=1}^H \epsilon_{\hat{\theta}}(x, y_{1:h-1}) \right\} \\
&\quad + 400\Delta^2\hat{\mathbb{E}}_{\mathcal{D}} \left[\sum_{h=1}^H \text{Var}_{\pi_D}(x, y_{1:h-1}) \right] + \frac{16A \log(4J/\delta)}{n} + 100\Delta^2\sigma_{\star}^2.
\end{aligned}$$

where the first inequality uses (60), the second inequality uses [Lemma K.2](#) (1), and the third inequality uses (61). Therefore, we denote $\sigma^2(\mathcal{D}) := \hat{\mathbb{E}}_{\mathcal{D}} \left[\sum_{h=1}^H \text{Var}_{\pi_D}(x, y_{1:h-1}) \right]$, and we have shown that for any $\delta \in (0, 1)$, any $\Delta(0, \frac{1}{100B}]$, it holds that

$$\mathbb{P}_{\mathcal{D} \sim \pi_D} \left(D_{\text{seq},N}(\pi_D \parallel \pi_{\hat{\theta}}) \leq C \left(E_1 + \Delta^2 \sigma^2(\mathcal{D}) + \Delta^2 \sigma_{\star}^2 + \frac{A}{n} \left(\frac{1}{\Delta^2} + \log(1/\delta) \right) \right) \right) \leq \delta,$$

where $C > 0$ is an absolute constant.

By the arbitrariness of $\delta \in (0, 1)$, taking expectation gives

$$\begin{aligned}
\mathbb{E} [D_{\text{seq},N}(\pi_D \parallel \pi_{\hat{\theta}})] &\leq C \left(\mathbb{E}[E_1] + \Delta^2 \mathbb{E}[\sigma^2(\mathcal{D})] + \Delta^2 \sigma_{\star}^2 + \frac{A}{n} \left(\frac{1}{\Delta^2} + 1 \right) \right) \\
&\leq 2C \left(\sqrt{\frac{\sigma_{\star}^2}{n}} + \Delta^2 \sigma_{\star}^2 + \frac{A}{n\Delta^2} \right).
\end{aligned}$$

Choosing $\Delta = \min \left\{ \frac{1}{100B}, \left(\frac{A}{\sigma_{\star}^2 n} \right)^{1/4} \right\}$ completes the proof. $\square$

K.3.1 PROOFS OF THE SUPPORTING LEMMAS

Proof of [Lemma K.1](#). Recall that $\hat{\pi} = \pi_{\hat{\theta}}$, where $\hat{\theta} = \arg \max_{\theta \in \Theta} \sum_{(x, y_{1:H}) \in \mathcal{D}} \log \pi_{\theta}(y_{1:H} \mid x)$. Then by the concavity, we know

$$\left\langle \hat{\mathbb{E}}_{\mathcal{D}} [\log \pi_{\theta}(y_{1:H} \mid x)], \theta - \hat{\theta} \right\rangle \leq 0, \quad \forall \theta \in \Theta$$

where we recall that $\hat{\mathbb{E}}_{\mathcal{D}}$ is the empirical distribution $(x, y_{1:H}) \sim \text{Unif}(\mathcal{D})$. Using the expressing (54) and $\theta^{\star} \in \Theta$, we know

$$\hat{\mathbb{E}}_{\mathcal{D}} \left[\sum_{h=1}^H \left\langle (\phi(x, y_{1:h}) - \bar{\phi}_{\hat{\theta}}(x, y_{1:h-1})), \theta^{\star} - \hat{\theta} \right\rangle \right] \leq 0.$$

Therefore, combining the inequality above with [Eq. \(56\)](#), we have

$$\begin{aligned}
\widehat{\mathbb{E}}_{\mathcal{D}} \left[\sum_{h=1}^H \epsilon_{\widehat{\theta}}(x, y_{1:h-1}) \right] &= \widehat{\mathbb{E}}_{\mathcal{D}} \left[\sum_{h=1}^H D_{\text{KL}}(\pi_{\mathcal{D}}(\cdot \mid x, y_{1:h-1}) \parallel \widehat{\pi}(\cdot \mid x, y_{1:h-1})) \right] \\
&\leq \widehat{\mathbb{E}}_{\mathcal{D}} \left[\sum_{h=1}^H \left\langle \bar{\phi}_{\theta^*}(x, y_{1:h-1}) - \bar{\phi}_{\widehat{\theta}}(x, y_{1:h-1}), \theta^* - \widehat{\theta} \right\rangle \right] \\
&\leq \widehat{\mathbb{E}}_{\mathcal{D}} \left[\sum_{h=1}^H \left\langle \bar{\phi}_{\theta^*}(x, y_{1:h-1}) - \phi(x, y_{1:h}), \theta^* - \widehat{\theta} \right\rangle \right] \\
&\leq 2 \left\| \widehat{\mathbb{E}}_{\mathcal{D}} \left[\sum_{h=1}^H \phi^*(x, y_{1:h}) \right] \right\| =: E'_1,
\end{aligned}$$

where we recall that $\phi^*(x, y_{1:h}) := \phi(x, y_{1:h}) - \bar{\phi}_{\theta^*}(x, y_{1:h-1})$. By definition, it holds that $\mathbb{E}_{\pi_0}[\phi^*(x, y_{1:h}) \mid x, y_{1:h-1}] = 0$, and hence

$$\begin{aligned}
\mathbb{E}(E'_1)^2 &= \mathbb{E} \left\| \widehat{\mathbb{E}}_{\mathcal{D}} \left[\sum_{h=1}^H \phi^*(x, y_{1:h}) \right] \right\|^2 \\
&= \frac{1}{n} \mathbb{E}_{\pi_0} \left\| \sum_{h=1}^H \phi^*(x, y_{1:h}) \right\|^2 = \frac{1}{n} \mathbb{E}_{\pi_0} \left[\sum_{h=1}^H \|\phi^*(x, y_{1:h})\|^2 \right] = \frac{\sigma_*^2}{n}.
\end{aligned}$$

This gives the desired upper bound. $\square$

Proof of [Lemma K.2](#). By Freedman inequality ([Lemma H.3](#)) and union bound, it is clear that (1) holds with probability at least $1 - \frac{\delta}{2}$. In the following, we prove (2).

Define the following weight function $\alpha = \alpha_{\widehat{\theta}} : \mathcal{X} \times \mathcal{A}^* \rightarrow [0, 1]$:

$$\alpha_{\widehat{\theta}}(x, y_{1:h-1}) = \begin{cases} 1, & \sum_{j \leq h-1} \epsilon_{\widehat{\theta}}(x, y_{1:j}) \leq A, \\ 0, & \sum_{j < h-1} \epsilon_{\widehat{\theta}}(x, y_{1:j}) \geq A, \\ \frac{A - \sum_{j < h-1} \epsilon_{\widehat{\theta}}(x, y_{1:j})}{\epsilon_{\widehat{\theta}}(x, y_{1:h-1})}, & \text{otherwise.} \end{cases}$$

Then, by [Lemma K.4](#), it holds that for any $\theta \in \Theta$,

$$\begin{aligned}
\mathbb{E}_{\pi_0} \min \left\{ A, \sum_{h=1}^H \epsilon_{\widehat{\theta}}(x, y_{1:h-1}) \right\} &\leq 2 \mathbb{E}_{\pi_0} \min \left\{ A, \sum_{h=1}^H \epsilon_{\theta}(x, y_{1:h-1}) \right\} \\
&\quad + 2 \mathbb{E}_{\pi_0} \left[\sum_{h=1}^H \alpha(x, y_{1:h-1}) F(\epsilon_{\widehat{\theta}}(x, y_{1:h-1}), \epsilon_{\theta}(x, y_{1:h-1})) \right],
\end{aligned}$$

and

$$\begin{aligned}
\widehat{\mathbb{E}}_{\mathcal{D}} \min \left\{ A, \sum_{h=1}^H \epsilon_{\theta}(x, y_{1:h-1}) \right\} &\leq 2 \widehat{\mathbb{E}}_{\mathcal{D}} \min \left\{ A, \sum_{h=1}^H \epsilon_{\widehat{\theta}}(x, y_{1:h-1}) \right\} \\
&\quad + \widehat{\mathbb{E}}_{\mathcal{D}} \left[\sum_{h=1}^H \alpha(x, y_{1:h-1}) F(\epsilon_{\widehat{\theta}}(x, y_{1:h-1}), \epsilon_{\theta}(x, y_{1:h-1})) \right],
\end{aligned}$$

Therefore, it remains to control the error $\sum_{h=1}^H \alpha(x, y_{1:h-1}) F(\epsilon_{\widehat{\theta}}(x, y_{1:h-1}), \epsilon_{\theta}(x, y_{1:h-1}))$ under both $\mathbb{E}_{\pi_0}[\cdot]$ and $\widehat{\mathbb{E}}_{\mathcal{D}}[\cdot]$. We prove the following lemma, which leverages the structure of Gaussian distribution.

Lemma K.3. For any $K \geq 1$, $\Delta \in (0, \frac{1}{40KB}]$, $\theta \in \mathbb{B}_2(1)$, distributions $\mu_1, \dots, \mu_K$ over $\mathcal{Z} := \mathcal{X} \times \mathcal{A}^*$, and weight function $\alpha : \mathcal{Z} \rightarrow [0, 1]$, it holds that

$$-\log \mathbb{P}_{\theta' \sim \mathcal{N}(0, \Delta^2)} (\forall i \in [K], \mathbb{E}_{z \sim \mu_i} \alpha(z) F(\epsilon_{\theta}(z), \epsilon_{\theta'}(z))) \leq 22K^2 \Delta^2 \mathbb{E}_{z \sim \mu_i} \text{Var}_{\pi_0}(z) \leq \frac{1}{\Delta^2} + 2.$$

In the following, we apply [Lemma K.3](#) with $K = 2$, parameter $\theta = \hat{\theta}$, weight function α , and the distributions μ_1, μ_2 defined as follows:

- Let μ_1 be the distribution of $x' = (x, y_{1:h-1})$ under $x \sim \mu, y_{1:H} \sim \pi_D(\cdot | x)$ and $h \sim \text{Unif}([H])$.
- Let μ_2 be the distribution of $x' = (x^t, y_{1:h-1}^t)$ under $t \sim \text{Unif}([n])$ and $h \sim \text{Unif}([H])$.

By definition, it holds that

$$\begin{aligned}\mathbb{E}_{z \sim \mu_1} \alpha(z) F(\epsilon_\theta(z), \epsilon_{\theta'}(z)) &= \frac{1}{H} \mathbb{E}_{\pi_D} \left[\sum_{h=1}^H \alpha(x, y_{1:h-1}) F(\epsilon_{\hat{\theta}}(x, y_{1:h-1}), \epsilon_\theta(x, y_{1:h-1})) \right], \\ \mathbb{E}_{z \sim \mu_1} \text{Var}_{\pi_D}(z) &= \frac{1}{H} \mathbb{E}_{\pi_D} \left[\sum_{h=1}^H \text{Var}_{\pi_D}(x, y_{1:h-1}) \right] = \frac{\sigma_\star^2}{H}, \\ \mathbb{E}_{z \sim \mu_2} \alpha(z) F(\epsilon_\theta(z), \epsilon_{\theta'}(z)) &= \frac{1}{H} \widehat{\mathbb{E}}_D \left[\sum_{h=1}^H \alpha(x, y_{1:h-1}) F(\epsilon_{\hat{\theta}}(x, y_{1:h-1}), \epsilon_\theta(x, y_{1:h-1})) \right], \\ \mathbb{E}_{z \sim \mu_2} \text{Var}_{\pi_D}(z) &= \frac{1}{H} \widehat{\mathbb{E}}_D \left[\sum_{h=1}^H \text{Var}_{\pi_D}(x, y_{1:h-1}) \right].\end{aligned}$$

Then, we consider the following set

$$\Theta_\theta^\pm := \{\forall i \in \{1, 2\}, \mathbb{E}_{z \sim \mu_i} \alpha(z) F(\epsilon_\theta(z), \epsilon_{\theta'}(z)) \leq 100\Delta^2 \mathbb{E}_{z \sim \mu_i} \text{Var}_{\pi_D}(z)\}.$$

By [Lemma K.3](#), it holds that

$$q_{\hat{\theta}} := \mathbb{P}_{\theta' \sim \mathcal{N}(0, \Delta^2 I)}(\theta' \in \Theta_\theta^\pm) \geq \exp\left(-\frac{1}{\Delta^2} - 2\right).$$

Therefore, we have

$$\begin{aligned}\mathbb{P}(\forall j \in [N], \theta_j \notin \Theta_\theta^\pm | \hat{\theta}) &= \mathbb{P}_{\theta_1, \dots, \theta_J \sim \mathcal{N}(0, \Delta^2 I)}(\forall j \in [N], \theta_j \notin \Theta_\theta^\pm) \\ &\leq (1 - q_{\hat{\theta}})^N \leq \exp(-Nq_{\hat{\theta}}) \leq \frac{\delta}{2},\end{aligned}$$

and hence $\mathbb{P}(\exists j \in [N], \theta_j \in \Theta_\theta^\pm) \geq 1 - \frac{\delta}{2}$. The proof of (2) is hence completed. $\square$

Proof of [Lemma K.3](#). By definition, we have $\pi_{\theta'}(y | z) \propto_y \pi_\theta(y | z) \cdot \exp(\langle \theta' - \theta, \phi(z, y) \rangle)$, i.e.,

$$\log \pi_{\theta'}(y | z) - \log \pi_\theta(y | z) = \langle \theta' - \theta, \phi(z, y) \rangle - \log \mathbb{E}_{y \sim \pi_\theta(\cdot | z)} \exp(\langle \theta' - \theta, \phi(z, y) \rangle).$$

Therefore,

$$\begin{aligned}\epsilon_\theta(z) - \epsilon_{\theta'}(z) &= D_{\text{KL}}(\pi_D(\cdot | z) \| \pi_\theta(\cdot | z)) - D_{\text{KL}}(\pi_D(\cdot | z) \| \pi_{\theta'}(\cdot | z)) \\ &= \mathbb{E}_{\pi_D(\cdot | z)} \langle \theta' - \theta, \phi(z, y) \rangle - \log \mathbb{E}_{y \sim \pi_\theta(\cdot | z)} \exp(\langle \theta' - \theta, \phi(z, y) \rangle) \\ &= \langle \theta' - \theta, \bar{\phi}_{\theta^\star}(z) - \bar{\phi}_\theta(z) \rangle - \log \mathbb{E}_{y \sim \pi_\theta(\cdot | z)} \exp(\langle \theta' - \theta, \phi(z, y) - \bar{\phi}_\theta(z) \rangle),\end{aligned}$$

where we recall that $\bar{\phi}_\theta(z) = \mathbb{E}_{y \sim \pi_\theta(\cdot | z)}[\phi(z, y)]$.

In the following, we denote $\phi_\theta(z, y) := \phi(z, y) - \bar{\phi}_\theta(z)$, and

$$\begin{aligned}E_{\theta'}^+(z) &:= \log \mathbb{E}_{y \sim \pi_\theta(\cdot | z)} \exp(\langle \theta' - \theta, \phi_\theta(z, y) \rangle), \\ E_{\theta'}^-(z) &:= \langle \theta' - \theta, \bar{\phi}_{\theta^\star}(z) - \bar{\phi}_\theta(z) \rangle.\end{aligned}$$

We first bound $E_{\theta'}^+(z)$. By definition, we have $E_{\theta'}^+(z) = D_{\text{KL}}(\pi_\theta(\cdot | z) \| \pi_{\theta'}(\cdot | z)) \geq 0$. Further, using Jensen's inequality, for any $z \in \mathcal{Z}$, we have

$$\begin{aligned}\mathbb{E}_{\theta' \sim \mathcal{N}(\theta, \Delta^2 I)}[E_{\theta'}^+(z)] &\leq \log \mathbb{E}_{\theta' \sim \mathcal{N}(\theta, \Delta^2 I)} \mathbb{E}_{y \sim \pi_\theta(\cdot | z)}[\exp(\langle \theta' - \theta, \phi_\theta(z, y) \rangle)] \\ &= \log \mathbb{E}_{y \sim \pi_\theta(\cdot | z)} \exp\left(\frac{1}{2} \Delta^2 \|\phi_\theta(z, y)\|^2\right) \\ &\leq \Delta^2 \mathbb{E}_{y \sim \pi_\theta(\cdot | z)} \|\phi_\theta(z, y)\|^2,\end{aligned}$$

where the last inequality follows from $e^t \leq 1 + 2t$ for $t \in [0, 1]$. Further, using [Lemma H.5](#), we have

$$\begin{aligned} \mathbb{E}_{y \sim \pi_\theta(\cdot|z)} \|\phi_\theta(z, y)\|^2 &= \mathbb{E}_{y \sim \pi_\theta(\cdot|z)} \|\phi(z, y) - \phi_\theta(z)\|^2 \\ &\leq 3 \mathbb{E}_{y \sim \pi_0(\cdot|z)} \|\phi(z, y) - \phi_{\theta^*}(z)\|^2 + 4B^2 D_{\text{KL}}(\pi_0(\cdot|z) \parallel \pi_\theta(\cdot|z)) \\ &= 3\text{Var}_{\pi_0}(z) + 4B^2 \epsilon_\theta(z). \end{aligned}$$

Next, we bound $|E_{\theta'}^-(z)|$. Under $\theta' \sim \mathcal{N}(\theta, \Delta^2 I)$, it is clear that $\langle \theta' - \theta, \bar{\phi}_{\theta^*}(z) - \bar{\phi}_\theta(z) \rangle \sim \mathcal{N}(0, \Delta^2 \|\bar{\phi}_{\theta^*}(z) - \bar{\phi}_\theta(z)\|^2)$ for any fixed z . Therefore, it holds that

$$\begin{aligned} \mathbb{E}_{\theta' \sim \mathcal{N}(\theta, \Delta^2 I)} |E_{\theta'}^-(z)| &= \sqrt{\frac{2}{\pi}} \Delta \cdot \|\bar{\phi}_{\theta^*}(z) - \bar{\phi}_\theta(z)\| \\ &\leq \Delta \cdot \left(2\sqrt{\text{Var}_{\pi_0}(z) \cdot \epsilon_\theta(z)} + 3B\epsilon_\theta(z) \right) \\ &\leq \left(\frac{1}{8K} + 3B\Delta \right) \epsilon_\theta(z) + 8K\Delta^2 \text{Var}_{\pi_0}(z). \end{aligned}$$

where the second line uses [\(56\)](#).

Combining the inequalities above, we know that for $i \in [K]$, it holds that

$$\begin{aligned} \mathbb{E}_{\theta' \sim \mathcal{N}(\theta, \Delta^2 I)} [\mathbb{E}_{z \sim \mu_i} [\alpha(z) E_{\theta'}^+(z)]] &\leq \Delta^2 \mathbb{E}_{z \sim \mu_i} [3\text{Var}_{\pi_0}(z) + 4B^2 \alpha(z) \epsilon_\theta(z)], \\ \mathbb{E}_{\theta' \sim \mathcal{N}(\theta, \Delta^2 I)} [\mathbb{E}_{z \sim \mu_i} [\alpha(z) |E_{\theta'}^-(z)|]] &\leq \mathbb{E}_{z \sim \mu_i} \left[8K\Delta^2 \text{Var}_{\pi_0}(z) + \left(\frac{1}{8K} + 3B\Delta \right) \alpha(z) \epsilon_\theta(z) \right], \end{aligned}$$

and hence by Markov's inequality, it holds that $p := \mathbb{P}_{\theta' \sim \mathcal{N}(\theta, \Delta^2 I)}(\theta' \notin \Theta^-) \geq \frac{1}{2}$, where we denote $\Theta^- = \cup_{i \in [K]} \Theta_i$, and

$$\Theta_i := \left\{ \theta' \in \mathbb{R}^d : \mathbb{E}_{z \sim \mu_i} \alpha(z) |\epsilon_\theta(z) - \epsilon_{\theta'}(z)| \geq \mathbb{E}_{z \sim \mu_i} \left[(6K + 16K^2) \Delta^2 \text{Var}_{\pi_0}(z) + \frac{1}{2} \alpha(z) \epsilon_\theta(z) \right] \right\}.$$

Note that $D_{\text{KL}}(\mathcal{N}(\theta, \Delta^2 I) \parallel \mathcal{N}(0, \Delta^2 I)) = \frac{\|\theta\|^2}{2\Delta^2} \leq \frac{1}{2\Delta^2}$. Hence, by data-processing inequality, we can bound $q := \mathbb{P}_{\theta' \sim \mathcal{N}(0, \Delta^2 I)}(\theta' \notin \Theta^-)$ as

$$\begin{aligned} \frac{1}{2\Delta^2} &\geq D_{\text{KL}}(\mathcal{N}(\theta, \Delta^2 I) \parallel \mathcal{N}(0, \Delta^2 I)) \geq D_{\text{KL}}(\text{Ber}(p) \parallel \text{Ber}(q)) \\ &= p \log \frac{p}{q} + (1-p) \log \frac{1-p}{1-q} \geq \frac{1}{2} \log(1/q) - \log 2, \end{aligned}$$

implying that $-\log q \leq \frac{1}{\Delta^2} + 2$. This is the desired result. $\square$

Lemma K.4. Suppose that $a_1, \dots, a_H, b_1, \dots, b_H \geq 0$. Let

$$\alpha_h = \begin{cases} 1, & \sum_{j \leq h} a_j \leq A, \\ 0, & \sum_{j \leq h} a_j > A, \\ \frac{A - \sum_{j < h} a_j}{a_h}, & \text{otherwise.} \end{cases}$$

Then clearly $\alpha_h \in [0, 1] \forall h \in [H]$, and it holds that $\sum_{h=1}^H \alpha_h a_h = \min \left\{ A, \sum_{h=1}^H a_h \right\}$, and

$$\begin{aligned} \min \left\{ A, \sum_{h=1}^H a_h \right\} &\leq 2 \min \left\{ A, \sum_{h=1}^H b_h \right\} + 2 \sum_{h=1}^H \alpha_h F(a_h, b_h), \\ \min \left\{ A, \sum_{h=1}^H b_h \right\} &\leq 2 \min \left\{ A, \sum_{h=1}^H a_h \right\} + \sum_{h=1}^H \alpha_h F(a_h, b_h), \end{aligned}$$

where we recall that $F(a, b) = |a - b| - \frac{1}{2}a$.

Proof of Lemma K.4. Fix the sequence $a_1, \dots, a_H$.

We first prove $\sum_{h=1}^H \alpha_h a_h = \min \left\{ A, \sum_{h=1}^H a_h \right\}$.

Case 1: $\sum_{h=1}^H a_h \leq A$. In this case, $\alpha_h = 1 \forall h \in [H]$, and the equation holds trivially.

Case 2: $\sum_{h=1}^H a_h > A$. In this case, we let $\ell \in [H]$ be the maximal index such that $\alpha_\ell > 0$. Then, by definition, $\sum_{j < \ell} a_j \leq A$ and $\sum_{j \leq \ell} a_j > A$, and $\alpha_\ell = \frac{A - \sum_{j < \ell} a_j}{a_\ell}$. Hence,

$$\sum_{h=1}^H \alpha_h a_h = \sum_{h=1}^{\ell} \alpha_h a_h = \sum_{j < \ell} a_j + \alpha_\ell a_\ell = A.$$

We note that from the proof above, we also know that for any sequence $(c_1, \dots, c_H)$ such that $c_h \geq a_h$ for $h \in [H]$, we have $\min \left\{ A, \sum_{h=1}^H c_h \right\} \leq \sum_{h=1}^H \alpha_h c_h$.

Next, we prove the inequalities. We note that

$$\sum_{h=1}^H \alpha_h F(a_h, b_h) = \sum_{h=1}^H \alpha_h |a_h - b_h| - \frac{1}{2} \sum_{h=1}^H \alpha_h a_h,$$

or equivalently,

$$\sum_{h=1}^H \alpha_h |a_h - b_h| = \sum_{h=1}^H \alpha_h F(a_h, b_h) + \frac{1}{2} \min \left\{ A, \sum_{h=1}^H a_h \right\}.$$

Therefore,

$$\begin{aligned} \min \left\{ A, \sum_{h=1}^H a_h \right\} &= \sum_{h=1}^H \alpha_h a_h \leq \min \left\{ A, \sum_{h=1}^H b_h \right\} + \sum_{h=1}^H \alpha_h |a_h - b_h| \\ &= \min \left\{ A, \sum_{h=1}^H b_h \right\} + \sum_{h=1}^H \alpha_h F(a_h, b_h) + \frac{1}{2} \min \left\{ A, \sum_{h=1}^H a_h \right\}. \end{aligned}$$

Re-organizing yields the first inequality. Similarly, we have

$$\begin{aligned} \min \left\{ A, \sum_{h=1}^H b_h \right\} &\leq \min \left\{ A, \sum_{h=1}^H (a_h + |a_h - b_h|) \right\} \leq \sum_{h=1}^H \alpha_h (a_h + |a_h - b_h|) \\ &= \frac{3}{2} \min \left\{ A, \sum_{h=1}^H a_h \right\} + \sum_{h=1}^H \alpha_h F(a_h, b_h). \end{aligned}$$

The proof is hence completed. $\square$

K.4 PROOF OF PROPOSITION 5.1 (LOWER BOUND)

In the following, we construct $\mathcal{X} = \mathbb{R} \sqcup \{-, +\}$, $\mathcal{Y} = \{-1, 0, 1\}$ and $\Theta = \mathbb{B}_2(1)$ with $d = 2$. The feature map ϕ satisfies $\phi(x, y_{1:h}) = \phi(x, y_h)$, i.e., $y_{1:H} \sim \pi_\theta(\cdot | x)$ are i.i.d. We fix $B \geq c_B \log H$ for a sufficiently large constant $c_B > 0$.

Case 1: $\eta \geq \frac{8}{HB}$. We define $\bar{\eta} := \eta \cdot HB$ and $\alpha = \frac{\bar{\eta}}{2(\bar{\eta}-1)} \leq \frac{5}{8}$. Let

$$v_0 = [1; 0], \quad v_1 = [\alpha; \sqrt{1-\alpha^2}], \quad v_{-1} = [\alpha; -\sqrt{1-\alpha^2}].$$

For $y_{1:h} \in \mathcal{A}^h$, we define $\phi(\eta, y_{1:h}) = Bv_{y_h}$, and we consider the problem instance with μ supported on $x = \eta$ (i.e., $\mu(\eta) = 1$) and $\theta^* = v_0$.

In the following, we omit the dependence on $x = \eta$. Then, under this construction, we have

$$\pi_\theta(y_h | y_{1:h}) = \frac{\exp(B\langle \theta, v_{y_h} \rangle)}{\sum_{y' \in \mathcal{Y}} \exp(B\langle \theta, v_{y'} \rangle)} = \pi_\theta(y_h).$$

We study the SGD update starting from $\theta^0 = v_1$. By definition,

$$\nabla \log \pi_\theta(y_{1:H}) = \sum_{h=1}^H \left(\phi(y_h) - \mathbb{E}_{y \sim \pi_\theta} [\phi(y)] \right).$$

In the following, we denote $\hat{F}(y_{1:H}) := \frac{1}{H} \sum_{h=1}^H v_{y_h}$, and

$$F(\theta) := \mathbb{E}_{y \sim \pi_\theta} [\phi(y)] = \frac{\sum_{y \in \mathcal{Y}} y \exp(B \langle \theta, v_y \rangle)}{\sum_{y' \in \mathcal{Y}} \exp(B \langle \theta, v_{y'} \rangle)}$$

Then, the SGD update can be written as

$$\theta^{t+1} = \text{Proj}_\Theta \left(\theta^t + \bar{\eta} \left(\hat{F}(y_{1:H}^t) - F(\theta^t) \right) \right).$$

We make the following claims.

Claim 1. For $y \in \mathcal{Y}$ and $\|\theta - v_y\| \leq \frac{1}{16}$, it holds that $1 - \pi_\theta(y) \leq 2e^{-B/4} =: \epsilon_1$ and hence $\|F(\theta) - v_y\| \leq 2\epsilon_1$.

Claim 2. Suppose that $\epsilon_1 \leq \min\{\frac{1}{4nH}, \frac{1}{5HB^2}\}$. Then it holds that $\text{Var}_{\pi_0}[\phi(y_1)] \leq \frac{1}{H}$ and $\sigma_* \leq 1$, and $\text{Cov}_N(\pi_0 \parallel \pi_\theta) \geq 1 - \frac{1}{2n}$ for $\log N \leq \frac{HB}{8}$ and $\theta \in \Theta$ such that $\min\{\|\theta - v_1\|, \|\theta - v_{-1}\|\} \leq \frac{1}{16}$. Further, with probability at least 0.5, it holds that $\hat{F}(y_{1:H}^t) = e_0$ for all $t \in [n]$.

In the following, we condition on this event.

Claim 3. By definition, for $y \in \{-1, 1\}$, we have $\|v_y + \bar{\eta}(v_0 - v_y)\| = \bar{\eta} - 1$ and $v_{-1} + v_1 = \frac{\bar{\eta}}{\bar{\eta}-1} v_0$.

Claim 4. Let $\epsilon = 16\epsilon_1 + 4\epsilon_0$. Suppose that $\epsilon \leq \frac{1}{16}$. Then if $\|\theta^t - v_y\| \leq \epsilon$, then it holds that $\|\theta^{t+1} - v_{-y}\| \leq \epsilon$.

Combining the above claims, we know that there is a constant C such that as long as $B \geq C \log(nH)$, it holds that $\sigma_* \leq 1$ and with probability at least 0.5, $\|\theta^t - v_{t \bmod 2}\| \leq \frac{1}{16}$ for all $t \in [n]$. Therefore, by Claim 2, this gives $\text{Cov}_N(\pi_{\theta^t}) \geq \frac{1}{2}$ as long as $\log N \leq \frac{HB}{8}$. $\square$

Proof of the claims. To prove Claim 1, we note that $\langle \theta, v_y \rangle \geq 1 - \|\theta - v_y\| \geq \frac{15}{16}$ and for $y' \neq y$, $\langle \theta, v_{y'} \rangle \leq \langle v_y, v_{y'} \rangle + \|\theta - v_y\| \leq \alpha + \frac{1}{16} \leq \frac{11}{16}$. Therefore,

$$1 - \pi_\theta(y) \leq \frac{\sum_{y' \neq y} e^{B \langle \theta, v_{y'} \rangle}}{e^{B \langle \theta, v_y \rangle}} \leq \frac{2}{e^{B/4}} = \epsilon_1.$$

In particular, we know $1 - \pi_0(0) \leq \epsilon_1$, and hence $\text{Var}_{\pi_0}[\phi(y_1)] \leq 5B^2\epsilon_1$. Further, we also know $\mathbb{P}_{\pi_0}(y_h = 0 \forall h \in [H]) \geq (1 - \epsilon_1)^H \geq 1 - H\epsilon_1$. Therefore, taking the union bound, we know $\mathbb{P}(y_h^t = 0 \forall h \in [H], t \in [n]) \geq 1 - nH\epsilon_1 \geq \frac{1}{2}$.

Furthermore, for any θ such that $\min\{\|\theta - v_1\|, \|\theta - v_{-1}\|\} \leq \frac{1}{16}$, as long as $\log N \leq H(\log(1 - \epsilon_1) - \log(\epsilon_1))$, we have

$$\text{Cov}_N(\pi_0 \parallel \pi_\theta) \geq \left(1 - \frac{1}{2n}\right) \mathbb{I}\{H \log \pi_0(0) - H \log \pi_\theta(0) \geq \log N\} \geq 1 - \frac{1}{2n}.$$

In particular, this is ensured when $\log N \leq \frac{HB}{8}$. This completes the proof of Claim 2.

Claim 3 follows immediately from the definition of α , v_0 , v_1 and v_{-1} . Finally, we prove claim 4. We define $u^t := \theta^t + \bar{\eta}(\hat{F}(y_{1:H}^t) - F(\theta^t))$. Then it holds that

$$\begin{aligned} \|u^t - (\bar{\eta} - 1)v_{-y}\| &= \|u^t - \bar{\eta}v_0 + (\bar{\eta} - 1)v_y\| \leq \|\theta^t - v_y\| + \bar{\eta}\|\hat{F}(y_{1:H}^t) - v_0\| + \bar{\eta}\|F(\theta^t) - v_y\| \\ &\leq \epsilon + \bar{\eta}(2\epsilon_1 + \epsilon_0) =: \epsilon'. \end{aligned}$$

In particular, it holds that $|\|u^t\| - (\bar{\eta} - 1)| \leq \epsilon'$ and hence $\|u^t\| \geq \bar{\eta} - 1 - \epsilon' \geq \bar{\eta} - 2 \geq 1$. Therefore, $\theta^{t+1} = \text{Proj}_\Theta(u^t) = \frac{u^t}{\|u^t\|}$, and we can bound

$$\begin{aligned} \|\theta^{t+1} - v_{-y}\| &= \left\| \frac{u^t - (\bar{\eta} - 1)v_{-y}}{\|u^t\|} + v_{-y} \left(\frac{\bar{\eta} - 1}{\|u^t\|} - 1 \right) \right\| \\ &\leq \frac{\|u^t - (\bar{\eta} - 1)v_{-y}\|}{\|u^t\|} + \frac{|\bar{\eta} - 1 - \|u^t\||}{\|u^t\|} \\ &\leq \frac{2\epsilon'}{\|u^t\|} \leq \frac{4\epsilon'}{\bar{\eta}} = \frac{4}{\bar{\eta}}\epsilon + 8\epsilon_1 + 4\epsilon_0 \leq \epsilon. \end{aligned}$$

□

Case 2: $\eta \leq \frac{8}{HB}$. Let $\bar{B} \leq B$ be a parameter such that $\bar{B} \geq c_B \log(c_B n H)$ for a sufficiently large constant c_B , and we again denote $\bar{\eta} = HB\bar{\eta}$.

In this case, we choose the distribution μ to be $\mu(+) = 1 - \mu(-) = \min\left\{1, \frac{BH}{512en\bar{B}^2 \log N}\right\}$, the feature map be specified as $\phi(-, \cdot) = 0$ and $\phi(+, y_{1:h}) = [y_h \bar{B}; 0]$ for $y \in \mathcal{Y}$. We choose $\theta^* = [1; 0]$. Note that $\pi_D(1 | +) = \frac{e^{\bar{B}}}{e^{-\bar{B}} + 1 + e^{\bar{B}}}$, and hence $1 - \pi_D(y_1 = 1 | +) \leq 2e^{-\bar{B}}$. Therefore, similar to Case 1, we have the following claims.

Claim 1. It holds that $\sigma_* \leq 1$, and with probability at least 0.5, it holds that $\sum_{t=1}^n \mathbb{I}\{x^t = +\} \leq 4\mu(1)n$, and for any t such that $x^t = +$, we have $y_h^t = 1$ for all $h \in [H]$.

In the following, we condition on this event.

Claim 2. For any $\theta \in \Theta$, it holds that $1 - \pi_\theta(1 | +) \leq \frac{2}{e^{\theta[1]\bar{B}}}$, and hence when $x^t = +$, we have

$$\|\nabla \log \pi_\theta(y^t | x^t)\| = \|H(\bar{B} - \mathbb{E}_{y_1 \sim \pi_\theta(\cdot | +)}[\phi(y_1)])\| \leq 2H\bar{B}|1 - \pi_\theta(1 | +)| \leq \frac{4H\bar{B}}{e^{\theta[1]\bar{B}}}.$$

Note that when $x^t = -$, we have $\nabla \log \pi_\theta(y^t | x^t) = 0$. Therefore, it holds that

$$0 \leq \theta^{t+1}[1] - \theta^t[1] \leq \mathbb{I}\{x^t = +\} \cdot \frac{4\bar{\eta}\bar{B}}{Be^{\theta^t[1]\bar{B}}}.$$

Claim 3. Suppose that $e^{\theta[1]\bar{B}} \leq \frac{H}{4 \log N}$. Then it holds that $\text{Cov}_N(\pi_\theta) \geq \frac{1}{2}$.

To complete the proof, we now choose $\theta' \in [-1, 1]$ such that $e^{\theta'\bar{B}} = \frac{H}{4 \log N}$, and we let $\theta^0 = [\theta' - \frac{1}{\bar{B}}; 0]$. Then, using Claim 2, we know that for any $t \in [n]$, it holds that

$$\theta^t[1] - \theta^0[1] \leq \sum_{t=1}^n \mathbb{I}\{x^t = +\} \cdot \frac{4\eta HB}{e^{\theta^0[1]\bar{B}}} \leq n \cdot \mu(+) \cdot \frac{16e\bar{\eta}\bar{B}}{Be^{\theta^0[1]\bar{B}}} \leq \mu(+) \cdot \frac{512e\bar{B}n \log N}{BH} \leq \frac{1}{\bar{B}}.$$

Therefore, we have $\theta^t[1] \leq \theta'$ and hence $\text{Cov}_N(\pi_{\theta^t}) \geq \frac{\mu(+)}{2}$ for any $t \in [n]$.

It remains to prove Claim 3. We note that similar to Claim 2, $\mathbb{P}_{\pi_D}(y_h = 1 \forall h \in [H] | x = +) \geq \frac{1}{2}$, and hence

$$\text{Cov}_N(\pi_\theta) \geq \frac{\mu(+)}{2} \cdot \mathbb{I}\{H(\log \pi_D(y_1 = 1 | +) - \log \pi_\theta(y_1 = 1 | +)) \geq \log N\} \geq \frac{\mu(+)}{2},$$

where we use $\log \pi_D(y_1 = 1 | +) \geq \log(1 - 2e^{-B}) \geq -3e^{-B}$ and $\log \pi_\theta(y_1 = 1 | +) \leq -\frac{1}{3e^{\theta[1]\bar{B}}}$. □

K.5 PROOF OF THE SUPPORTING RESULTS

We generalize Proposition 3.2 to show that in the worst case (where $\sigma_*^2 \asymp HB^2$), the scaling $\text{Cov}_N(\hat{\pi}) = \Omega(\frac{H}{n \log N})$ can be unavoidable for autoregressive linear model. This implies that the dependence on σ_*^2 is generally necessary to achieve upper bounds that do not explicitly scale with H .

Proposition K.1. Let $H, B, N, n \geq 1$, and assume $\log N \leq c \min\{H, B^2\}$ for a sufficiently small constant $c > 0$. There exists an instance of H -dimensional autoregressive linear model class Π with $\phi : \mathcal{X} \times \mathcal{A}^* \rightarrow \mathbb{B}_2(B)$ and $\Theta = \mathbb{B}_2(1)$, such that for any proper algorithm Alg with output $\hat{\pi} = \pi_{\hat{\theta}}$, there exists $\pi_D \in \Pi$, such that under π_D , it holds that

$$\mathbb{E}^{\pi_D, \text{Alg}}[\text{Cov}_N(\pi_D \parallel \hat{\pi})] \geq c \cdot \min\left\{1, \frac{H}{n \cdot \log N}\right\}.$$

Proof. We consider $\mathcal{X} = \{+, -\}$, $\mathcal{A} = \{0, 1\}$, and the distribution μ be given by $\mu(+) = 1 - \mu(-) = p$, where $p \in [0, 1]$ is a pre-specified parameter. Let the feature map ϕ be given by $\phi(y_{1:h} \mid -) = 0$, $\phi(y_{1:h} \mid +) = B y_h e_h$, where $(e_1, \dots, e_H)$ is a fixed orthonormal basis of $\mathbb{R}^H$. Note that with this construction, we have $\pi_\theta(y_h = \cdot \mid -, y_{1:h-1}) = \text{Ber}(1/2)$, and

$$\pi_\theta(y_h = \cdot \mid +, y_{1:h-1}) = \text{Ber}\left(\frac{e^{B\theta_h}}{1 + e^{B\theta_h}}\right) =: \pi_{\theta,h}.$$

Note that for any $h \in [H]$, we can bound

$$C_0 B |\theta_h - \theta'_h| \leq D_H(\pi_{\theta,h}, \pi_{\theta',h}) \leq C_1 B |\theta_h - \theta'_h|,$$

as long as $\theta_h \in [-\frac{1}{B}, \frac{1}{B}]$.

We fix $\epsilon \in [0, \frac{1}{\max\{\sqrt{H}, B\}}]$ to be determined later, and for any $v \in \{-1, 1\}^H$, we let $\theta_v := \epsilon \sum_{h=1}^H v_h e_h$, and

$$\Theta_0 := \{\theta_v : v \in \{-1, 1\}^H\} \subset \mathbb{B}_2(1), \quad \Pi_0 := \{\pi_\theta : \theta \in \Theta_0\}.$$

Then a direct argument shows that when $pn \leq \frac{c_0}{B^2 \epsilon^2}$ for a sufficiently small constant c_0 , there exists $\theta^* \in \Theta_0$ such that under $\pi_D = \pi_{\theta^*}$, it holds that

$$\sum_{h=1}^H \mathbb{P}^{\pi_D, \text{Alg}}(|\hat{\theta}_h - \theta_h^*| \geq \epsilon) \geq cH.$$

Therefore, with probability at least $\frac{c}{2}$, it holds that $\sum_{h=1}^H \mathbb{I}\{|\hat{\theta}_h - \theta_h^*| \geq \epsilon\} \geq \frac{cH}{2}$, and this in turn implies

$$\sum_{h=1}^H D_H^2(\pi_{\theta^*,h}, \pi_{\hat{\theta},h}) \geq c_1 H B^2 \epsilon^2.$$

Then, by [Proposition D.11](#), we know that under the above event, as long as $\log N \leq \frac{c_1 H B^2 \epsilon^2}{2}$, we have $\text{Cov}_N(\hat{\pi}) \geq \frac{p}{2}$. Choosing $\epsilon = \sqrt{\frac{4 \log N}{c_1 H B^2}}$ and $p = \min\{1, \frac{c_0}{n B^2 \epsilon^2}\}$ gives the desired lower bound. $\square$

L PROOFS FROM SECTION 6

L.1 PROOF OF THEOREM 6.1

Recall (from [Eq. \(11\)](#)) that we consider the token-level SGD iterates defined as

$$\theta^{t,h+1} = \text{Proj}_\Theta(\theta^{t,h} + \eta \nabla \log \pi_{\theta^{t,h}}(y_h^t \mid x^t, y_{1:h-1}^t)), \text{ for } h = 0, \dots, H-1, \quad (62)$$

and $\theta^{t+1} \equiv \theta^{t+1,0} := \theta^{t,H}$ for $t \in [T]$, where $(x^t, y_{1:H}^t) \sim \pi_D$.

To define the guarantee on θ^t which we are able to derive, we next define the following *test-time parameter update* $\vartheta^{\text{TTT}}(x, y_{1:h}; \theta)$, for a parameter θ and prompt x . It is defined recursively for $h = 0, 1, \dots, H-1$:

$$\vartheta^{\text{TTT}}(x, y_{1:h}; \theta) := \text{Proj}_\Theta(\vartheta^{\text{TTT}}(x, y_{1:h-1}; \theta) + \eta \nabla \log \pi_{\vartheta^{\text{TTT}}(x, y_{1:h-1}; \theta)}(y_h \mid x, y_{1:h-1})). \quad (63)$$

We then define a distribution $\pi_\theta^{\text{TTT}} : \mathcal{X} \rightarrow \Delta(\mathcal{Y}^H)$ as

$$\pi_\theta^{\text{TTT}}(\cdot \mid x, y_{1:h-1}) := \pi_{\vartheta^{\text{TTT}}(x, y_{1:h-1}; \theta)}(\cdot \mid x, y_{1:h-1}). \quad (64)$$

The distribution π_θ^{TTT} can be interpreted as an augmented version of the autoregressive linear model π_θ that performs test-time training during sampling.

Proof. We closely follow the proof of [Proposition 5.1](#) (cf. [Appendix K.1](#)).

We first note that by the proof of [Eq. \(57\)](#), we have

$$\mathbb{E} \left[\sum_{t=1}^T \langle -\nabla \log \pi_{\theta^{t,h}}(y_h^t | x^t, y_{1:h-1}^t), \theta^{t,h} - \theta^* \rangle \right] \leq \frac{1}{2\eta} + \eta \mathbb{E} \left[\sum_{t=1}^T \|\nabla \log \pi_{\theta^{t,h}}(y_h^t | x^t, y_{1:h-1}^t)\|^2 \right].$$

Further, by the proof of [Eq. \(58\)](#), we have

$$\begin{aligned} & \|\nabla \log \pi_{\theta^{t,h}}(y_h^t | x^t, y_{1:h-1}^t)\|^2 \\ & \leq 2 \|\nabla \log \pi_{\theta^*}(y_h^t | x^t, y_{1:h-1}^t)\|^2 \\ & \quad + 2B^2 \langle \nabla \log \pi_{\theta^*}(y_h^t | x^t, y_{1:h-1}^t) - \nabla \log \pi_{\theta^{t,h}}(y_h^t | x^t, y_{1:h-1}^t), \theta^{t,h} - \theta^* \rangle \end{aligned}$$

Note that the conditional distribution of $y_h^t | (x^t, y_{1:h-1}^t, \theta^{t,h})$ is given by $y_h^t \sim \pi_D(\cdot | x^t, y_{1:h-1}^t)$. Hence, taking expectation, we have

$$\begin{aligned} \mathbb{E} \left[\|\nabla \log \pi_{\theta^{t,h}}(y_h^t | x^t, y_{1:h-1}^t)\|^2 \right] & \leq 2 \mathbb{E}_{\pi_D} \|\nabla \log \pi_{\theta^*}(y_h | x, y_{1:h-1})\|^2 \\ & \quad + 2B^2 \mathbb{E} [\langle -\nabla \log \pi_{\theta^{t,h}}(y_h^t | x^t, y_{1:h-1}^t), \theta^{t,h} - \theta^* \rangle], \end{aligned}$$

and we also have (cf. [Eq. \(55\)](#))

$$\mathbb{E} [\langle -\nabla \log \pi_{\theta^{t,h}}(y_h^t | x^t, y_{1:h-1}^t), \theta^{t,h} - \theta^* \rangle] \geq \mathbb{E} D_{\text{KL}}(\pi_D(\cdot | x^t, y_{1:h-1}^t) \| \pi_{\theta^{t,h}}(\cdot | x^t, y_{1:h-1}^t)).$$

Combining the inequalities above, as long as $\eta \leq \frac{1}{4B^2}$, it holds that

$$\mathbb{E} \left[\sum_{t=1}^H \sum_{h=1}^H D_{\text{KL}}(\pi_D(\cdot | x^t, y_{1:h-1}^t) \| \pi_{\theta^{t,h}}(\cdot | x^t, y_{1:h-1}^t)) \right] \leq \frac{1}{\eta} + 4\eta T \sigma_*^2. \quad (65)$$

Finally, we note that

$$\theta^{t,h} = \vartheta^{\text{TTT}}(x^t, y_{h-1}^t; \theta^t),$$

and $x^t, y_{h-1}^t | \theta^t \sim \pi_D$. Therefore,

$$\begin{aligned} & \mathbb{E} [D_{\text{KL}}(\pi_D(\cdot | x^t, y_{1:h-1}^t) \| \pi_{\theta^{t,h}}(\cdot | x^t, y_{1:h-1}^t)) | \theta^t] \\ & = \mathbb{E}_{(x,y) \sim \pi_D} D_{\text{KL}}(\pi_D(\cdot | x, y_{1:h-1}) \| \pi_{\vartheta^{\text{TTT}}(x, y_{1:h-1}; \theta^t)}(\cdot | x, y_{1:h-1})) \\ & = \mathbb{E}_{(x,y) \sim \pi_D} D_{\text{KL}}(\pi_D(\cdot | x, y_{1:h-1}) \| \pi_{\theta^t}^{\text{TTT}}(\cdot | x, y_{1:h-1})). \end{aligned}$$

This implies

$$\begin{aligned} \frac{1}{\eta} + 4\eta T \sigma_*^2 & \geq \mathbb{E} \left[\sum_{t=1}^T \sum_{h=1}^H D_{\text{KL}}(\pi_D(\cdot | x^t, y_{1:h-1}^t) \| \pi_{\theta^{t,h}}(\cdot | x^t, y_{1:h-1}^t)) \right] \\ & = \mathbb{E} \left[\sum_{t=1}^T \mathbb{E} \left[\sum_{h=1}^H D_{\text{KL}}(\pi_D(\cdot | x^t, y_{1:h-1}^t) \| \pi_{\theta^{t,h}}(\cdot | x^t, y_{1:h-1}^t)) \mid \theta^t \right] \right] \\ & = \mathbb{E} \left[\sum_{t=1}^T \mathbb{E}_{\pi_D} \left[\sum_{h=1}^H D_{\text{KL}}(\pi_D(\cdot | x, y_{1:h-1}) \| \pi_{\theta^t}^{\text{TTT}}(\cdot | x, y_{1:h-1})) \right] \right] \\ & = \mathbb{E} \left[\sum_{t=1}^T D_{\text{KL}}(\pi_D \| \pi_{\theta^t}^{\text{TTT}}) \right], \end{aligned}$$

where the last equality uses the chain rule of KL divergence. $\square$

L.2 PROOF OF THEOREM F.2

We first note that by the proof of Eq. (57), we have

$$\mathbb{E} \left[\sum_{t=1}^T \langle \mathbb{E}_{(x,y) \sim \pi_0} [\hat{g}_{\theta^t}(y | x)], \theta^t - \theta^* \rangle \right] \leq \frac{1}{2\eta} + \eta \mathbb{E} \left[\sum_{t=1}^T \mathbb{E}_{(x,y) \sim \pi_0} \|\hat{g}_{\theta^t}(y | x)\|^2 \right].$$

In the following, we analyze $\langle \mathbb{E}_{(x,y) \sim \pi_0} [\hat{g}_{\theta}(y | x)], \theta^t - \theta^* \rangle$ and $\mathbb{E}_{(x,y) \sim \pi_0} \|\hat{g}_{\theta^t}(y | x)\|^2$ for any $\theta \in \Theta$, following the proof of Proposition 5.1 (cf. Appendix K.1).

We adopt the notation of Appendix K: For any θ and any pair $(x, y_{1:h-1})$, we denote $\bar{\phi}_{\theta}(x, y_{1:h-1}) = \mathbb{E}_{\pi_{\theta}}[\phi(x, y_{1:h}) | x, y_{1:h-1}]$ and

$$\epsilon_{\theta}(x, y_{1:h-1}) = D_{\text{KL}}(\pi_0(\cdot | x, y_{1:h-1}) \| \pi_{\theta}(\cdot | x, y_{1:h-1})).$$

By definition, we have (cf. Lemma K.4)

$$\sum_{h=1}^H \alpha_{\theta}(x, y_{1:h-1}) \epsilon_{\theta}(x, y_{1:h-1}) = \min \left\{ A, \sum_{h=1}^H \epsilon_{\theta}(x, y_{1:h-1}) \right\}, \quad (66)$$

and hence

$$\begin{aligned} \mathbb{E}_{(x,y) \sim \pi_0} \left[\sum_{h=1}^H \alpha_{\theta}(x, y_{1:h-1}) \epsilon_{\theta}(x, y_{1:h-1}) \right] &= \mathbb{E}_{(x,y) \sim \pi_0} \min \left\{ A, \sum_{h=1}^H \epsilon_{\theta}(x, y_{1:h-1}) \right\} \\ &= D_{\text{seq},N}(\pi_0 \| \pi_{\theta}), \end{aligned} \quad (67)$$

where we recall that $D_{\text{seq},N}(\pi_0 \| \pi_{\theta})$ is defined in Proposition D.10 and we denote $A = \log N$. Hence, by convexity (Eq. (55)),

$$\begin{aligned} &\langle \mathbb{E}_{(x,y) \sim \pi_0} [\hat{g}_{\theta}(y | x)], \theta - \theta^* \rangle \\ &= \mathbb{E}_{(x,y) \sim \pi_0} \left[\sum_{h=1}^H \alpha_{\theta}(x, y_{1:h-1}) \langle \bar{\phi}_{\theta^*}(x, y_{1:h-1}) - \bar{\phi}_{\theta}(x, y_{1:h-1}), \theta - \theta^* \rangle \right] \\ &\geq \mathbb{E}_{(x,y) \sim \pi_0} \left[\sum_{h=1}^H \alpha_{\theta}(x, y_{1:h-1}) \epsilon_{\theta}(x, y_{1:h-1}) \right] = D_{\text{seq},N}(\pi_0 \| \pi_{\theta}). \end{aligned}$$

Further, by Eq. (56), it holds that

$$\begin{aligned} &\|\hat{g}_{\theta}(y | x) - \hat{g}_{\theta^*}(y | x)\| \\ &\leq \sum_{h=1}^H \alpha_{\theta}(x, y_{1:h-1}) \|\bar{\phi}_{\theta^*}(x, y_{1:h-1}) - \bar{\phi}_{\theta}(x, y_{1:h-1})\| \\ &\leq \sum_{h=1}^H \alpha_{\theta}(x, y_{1:h-1}) \left(2\sqrt{\text{Var}_{\pi_0}(x, y_{1:h-1}) \cdot \epsilon_{\theta}(x, y_{1:h-1})} + 3B\epsilon_{\theta}(x, y_{1:h-1}) \right). \end{aligned}$$

Hence, using Eq. (66), we have

$$\begin{aligned} &\|\hat{g}_{\theta}(y | x) - \hat{g}_{\theta^*}(y | x)\|^2 \\ &\leq 8 \left(\sum_{h=1}^H \alpha_{\theta}(x, y_{1:h-1}) \sqrt{\text{Var}_{\pi_0}(x, y_{1:h-1}) \cdot \epsilon_{\theta}(x, y_{1:h-1})} \right)^2 \\ &\quad + 18B^2 \left(\sum_{h=1}^H \alpha_{\theta}(x, y_{1:h-1}) \epsilon_{\theta}(x, y_{1:h-1}) \right)^2 \\ &\leq 8 \left(\sum_{h=1}^H \alpha_{\theta}(x, y_{1:h-1}) \text{Var}_{\pi_0}(x, y_{1:h-1}) \right) \left(\sum_{h=1}^H \alpha_{\theta}(x, y_{1:h-1}) \epsilon_{\theta}(x, y_{1:h-1}) \right) \\ &\quad + 18AB^2 \left(\sum_{h=1}^H \alpha_{\theta}(x, y_{1:h-1}) \epsilon_{\theta}(x, y_{1:h-1}) \right) \\ &\leq 8A \left(\sum_{h=1}^H \text{Var}_{\pi_0}(x, y_{1:h-1}) \right) + 18AB^2 \left(\sum_{h=1}^H \alpha_{\theta}(x, y_{1:h-1}) \epsilon_{\theta}(x, y_{1:h-1}) \right). \end{aligned}$$

Therefore, taking expectation of $(x, y) \sim \pi_D$ and using $\mathbb{E}_{\pi_D} \|\hat{g}_{\theta^*}(y | x)\|^2 \leq \sigma_*^2$ and Eq. (67), it holds that

$$\mathbb{E}_{(x,y) \sim \pi_D} \|\hat{g}_{\theta}(y | x)\|^2 \leq (16A + 1)\sigma_*^2 + 36AB^2 D_{\text{seq},N}(\pi_D \| \pi_{\theta}), \quad \forall \theta.$$

Finally, combining the inequalities above, we know that

$$\begin{aligned} \mathbb{E} \left[\sum_{t=1}^T D_{\text{seq},N}(\pi_D \| \pi_{\theta^t}) \right] &\leq \mathbb{E} \left[\sum_{t=1}^T \langle \mathbb{E}_{(x,y) \sim \pi_D} [\hat{g}_{\theta^t}(y | x)], \theta^t - \theta^* \rangle \right] \\ &\leq \frac{1}{2\eta} + \eta \mathbb{E} \left[\sum_{t=1}^T \mathbb{E}_{(x,y) \sim \pi_D} \|\hat{g}_{\theta^t}(y | x)\|^2 \right] \\ &\leq \frac{1}{2\eta} + \eta T(16A + 1)\sigma_*^2 + 36AB^2 \mathbb{E} \left[\sum_{t=1}^T D_{\text{seq},N}(\pi_D \| \pi_{\theta^t}) \right]. \end{aligned}$$

Therefore, as long as $\eta \leq \frac{1}{72AB^2}$, it holds that

$$\mathbb{E} \left[\sum_{t=1}^T D_{\text{seq},N}(\pi_D \| \pi_{\theta^t}) \right] \leq \frac{1}{\eta} + 36\eta T A \sigma_*^2.$$

Optimally choosing η gives

$$\mathbb{E} \left[\frac{1}{T} \sum_{t=1}^T D_{\text{seq},N}(\pi_D \| \pi_{\theta^t}) \right] \lesssim \sqrt{\frac{\sigma_*^2 \log N}{T}} + \frac{B^2 \log N}{T}.$$

By Proposition D.10, this implies

$$\mathbb{E} \left[\frac{1}{T} \sum_{t=1}^T \text{Cov}_N(\pi_{\theta^t}) \right] \lesssim \sqrt{\frac{\sigma_*^2}{T \log N}} + \frac{B^2}{T}.$$

L.3 PROOFS FROM SECTION 6.2 (SELECTION)

Below we state and prove a generalization of Theorem 6.2 which holds when the data distribution π_D is not necessarily in the model class Π .

Theorem 6.2' (General version of Theorem 6.2). *Fix $N \geq 1$, and consider the estimator $\hat{\pi}$ from Eq. (13):*

$$\hat{\pi} := \arg \min_{\pi \in \Pi} \max_{\pi' \in \Pi} \widehat{\text{Cov}}_N(\pi' \| \pi). \quad (68)$$

For any $\delta \in (0, 1)$, parameter $a, c \geq 0$, with probability at least $1 - \delta$, it holds that

$$\text{Cov}_{N^{1+a+2c}}(\hat{\pi}) \lesssim \min_{\bar{\pi} \in \Pi} \text{Cov}_{N^a}(\bar{\pi}) + \frac{1}{N^{1-a-2c}} + \frac{\log \mathcal{N}_{\infty}(\Pi; c \log N) + \log \delta^{-1}}{n}. \quad (69)$$

Proof of Theorem 6.2'. Fix any $\bar{\pi} \in \Pi$ and $M, \alpha > 0$, and we study the estimator

$$\hat{\pi} := \arg \min_{\pi \in \Pi} \max_{\pi' \in \Pi} \widehat{\text{Cov}}_M(\pi' \| \pi). \quad (70)$$

By Lemma J.2, with probability at least $1 - \delta$, it holds that for $\forall \pi \in \Pi$,

$$\widehat{\text{Cov}}_N(\bar{\pi} \| \pi) \geq \frac{1}{2} \text{Cov}_{e^{2\alpha} N}^{\pi_D}(\bar{\pi} \| \pi) - \varepsilon_{\text{stat}},$$

where $\varepsilon_{\text{stat}} = \log(\mathcal{N}_{\infty}(\Pi; \alpha)/\delta)$. Next, again by Lemma J.2, with probability at least $1 - \delta$, it holds that for $\forall \pi \in \Pi$,

$$\widehat{\text{Cov}}_N(\pi \| \bar{\pi}) \leq 2 \text{Cov}_{e^{-2\alpha} N}^{\pi_D}(\pi \| \bar{\pi}) + \varepsilon_{\text{stat}}.$$

Therefore, we have with probability at least $1 - 2\delta$,

$$\begin{aligned} \frac{1}{2} \text{Cov}_{e^{2\alpha} N}^{\pi_D}(\bar{\pi} \| \hat{\pi}) - \varepsilon_{\text{stat}} &\leq \widehat{\text{Cov}}_N(\bar{\pi} \| \hat{\pi}) \leq \max_{\pi' \in \Pi} \widehat{\text{Cov}}_N(\pi' \| \hat{\pi}) \\ &= \min_{\pi \in \Pi} \max_{\pi' \in \Pi} \widehat{\text{Cov}}_N(\pi' \| \pi) \leq \max_{\pi' \in \Pi} \widehat{\text{Cov}}_N(\pi' \| \bar{\pi}) \\ &\leq 2 \max_{\pi' \in \Pi} \text{Cov}_{e^{-2\alpha} N}^{\pi_D}(\pi' \| \bar{\pi}) + \varepsilon_{\text{stat}}. \end{aligned}$$

Reorganizing yields

$$\text{Cov}_{e^{2\alpha}N}^{\pi_0}(\bar{\pi} \parallel \hat{\pi}) \leq 4 \max_{\pi \in \Pi} \text{Cov}_{e^{-2\alpha}N}^{\pi_0}(\pi \parallel \bar{\pi}) + 4\varepsilon_{\text{stat}}.$$

Note that for any N', N'' and policy π, π', π'' ,

$$\text{Cov}_{N'N''}^{\pi_0}(\pi' \parallel \pi) \leq \text{Cov}_{N'}^{\pi_0}(\pi' \parallel \pi'') + \text{Cov}_{N''}^{\pi_0}(\pi'' \parallel \pi). \quad (71)$$

Hence, for all π ,

$$\begin{aligned} \text{Cov}_{e^{2\alpha}N}^{\pi_0}(\pi_D \parallel \pi) &\leq \text{Cov}_{N'}^{\pi_0}(\pi_D \parallel \bar{\pi}) + \text{Cov}_{e^{2\alpha}N}^{\pi_0}(\bar{\pi} \parallel p^i), \\ \text{Cov}_{e^{-2\alpha}N}^{\pi_0}(\pi \parallel \bar{\pi}) &\leq \text{Cov}_{N'}^{\pi_0}(\pi \parallel \pi_D) + \text{Cov}_{e^{-2\alpha}N/N'}^{\pi_0}(\pi_D \parallel \bar{\pi}) \end{aligned}$$

Therefore, using the fact that $\text{Cov}_A^{\pi_0}(\pi \parallel \pi_D) \leq \frac{1}{A}$ and the inequalities above, we see that

$$\begin{aligned} \text{Cov}_{e^{2\alpha}N}^{\pi_0}(\hat{\pi}) &= \text{Cov}_{e^{2\alpha}N}^{\pi_0}(\pi_D \parallel \hat{\pi}) \\ &\leq \text{Cov}_{N'}^{\pi_0}(\pi_D \parallel \bar{\pi}) + \text{Cov}_{e^{2\alpha}N}^{\pi_0}(\bar{\pi} \parallel \hat{\pi}) \\ &\leq \text{Cov}_{N'}^{\pi_0}(\pi_D \parallel \bar{\pi}) + 4 \max_{\pi \in \Pi} \text{Cov}_{e^{-2\alpha}N}^{\pi_0}(\pi \parallel \bar{\pi}) + 4\varepsilon_{\text{stat}} \\ &\leq 5\text{Cov}_{N'}^{\pi_0}(\pi_D \parallel \bar{\pi}) + 4 \max_{\pi \in \Pi} \text{Cov}_{e^{-2\alpha}N/N'}^{\pi_0}(\pi \parallel \pi_D) + 4\varepsilon_{\text{stat}} \\ &\leq 5\text{Cov}_{N'}^{\pi_0}(\pi_D \parallel \bar{\pi}) + \frac{e^{2\alpha}N'}{N} + 4\varepsilon_{\text{stat}}. \end{aligned}$$

The claimed bound follows by setting $\bar{\pi} = \arg \min_{\pi \in \Pi} \text{Cov}_{N'}^{\pi_0}(\pi_D \parallel \pi)$, $\alpha = c \log N$, and $N' = N^a$. $\square$

L.4 PROOF OF THEOREM F.3

Divergence. For any distribution $P, Q \in \Delta(\mathcal{Y})$, we define the following divergence for $M \geq 1$:

$$\mathcal{E}_M(P \parallel Q) := \max \left\{ \mathbb{E}_{y \sim P} \left(\frac{dQ}{dP} - M \right)_+, \mathbb{E}_{y \sim Q} \left(\frac{dP}{dQ} - M \right)_+ \right\}.$$

Then, for policies $\pi, \pi' : \mathcal{X} \rightarrow \Delta(\mathcal{Y})$, we further define

$$\mathcal{E}_{M,\mu}(\pi \parallel \pi') := \mathbb{E}_{x \sim \mu} \mathcal{E}_M(\pi(\cdot \mid x) \parallel \pi'(\cdot \mid x)).$$

Under this divergence, it holds that for any event E ,

$$\mathbb{P}_{\mu,\pi}(E) \leq M \cdot \mathbb{P}_{\mu,\pi'}(E) + \mathcal{E}_{M,\mu}(\pi \parallel \pi'), \quad (72)$$

$$\mathbb{P}_{\mu,\pi'}(E) \leq M \cdot \mathbb{P}_{\mu,\pi}(E) + \mathcal{E}_{M,\mu}(\pi \parallel \pi'), \quad (73)$$

where $\mathbb{P}_{\mu,\pi}$ is the probability under $x \sim \mu$ and $y \sim \pi(\cdot \mid x)$. Furthermore, we can bound

$$\text{Cov}_{2M}(\pi) = \mathbb{P}_{\mu,\pi_0} \left(\frac{\pi_D(y \mid x)}{\pi(y \mid x)} \geq 2M \right) \leq \mathcal{E}_{M,\mu}(\pi_D \parallel \pi). \quad (74)$$

Theorem F.3' (General version of Theorem F.3). *Fix $N, \gamma \geq 1$ such that $N \geq 4\gamma^2$. Consider the estimator*

$$\hat{\pi} := \arg \min_{\pi \in \Pi} \max_{\pi' \in \Pi} \{ \widehat{\text{Cov}}_N(\pi' \parallel \pi) - 2\gamma \cdot \widehat{\text{Cov}}_N^{\pi}(\pi' \parallel \pi) \}. \quad (75)$$

Then with probability $1 - \delta$, it holds that

$$\text{Cov}_{2N\gamma}(\hat{\pi}) \lesssim \min_{\pi \in \Pi} \mathcal{E}_\gamma(\pi_D \parallel \pi) + \frac{\log(|\Pi|/\delta)}{n}.$$

Theorem F.3 is an immediate corollary by setting $\gamma = N^c$.

Proof of Theorem F.3'. For $\pi, \pi' \in \Pi$, we define the set

$$\mathcal{C}_N(\pi, \pi') = \left\{ (x, y) \mid \frac{\pi(y \mid x)}{\pi'(y \mid x)} \geq N \right\}.$$

Suppose an i.i.d. dataset $\mathcal{D} = \{(x^i, y^i)\}_{i \in [n]} \sim \pi_D$ is drawn. We write $\mu_n = \frac{1}{n} \sum_{i=1}^n \delta_{(x^i, y^i)}$ to denote the empirical measure (i.e., μ_n is the uniform distribution over $\mathcal{D}$). Note that

$$\widehat{\text{Cov}}_N(\pi' \parallel \pi) = \mu_n(\mathcal{C}_N(\pi', \pi)). \quad (76)$$

We also recall that

$$\widehat{\text{Cov}}_N^{\bar{\pi}}(\pi' \parallel \pi) := \frac{1}{n} \sum_{i=1}^n \mathbb{P}_{y \sim \bar{\pi}(\cdot | x^i)} \left(\frac{\pi'(y | x^i)}{\pi(y | x^i)} \geq N \right).$$

Therefore, we write $\widehat{\text{Cov}}_N^{\bar{\pi}}(\pi' \parallel \pi) = \mathbb{P}_{n, \bar{\pi}}(\mathcal{C}_N(\pi', \pi))$, where $\mathbb{P}_{n, \bar{\pi}}$ is the probability under the distribution $x \sim \mu_n, y \sim \bar{\pi}(\cdot | x)$.

Thus, the tournament estimator in Eq. (38) can be expressed as

$$\hat{\pi} := \arg \min_{\pi \in \Pi} \max_{\pi' \in \Pi} \mathcal{L}(\pi, \pi'), \quad (77)$$

where

$$\mathcal{L}(\pi, \pi') := \mu_n(\mathcal{C}_N(\pi', \pi)) - 2\gamma \cdot \mathbb{P}_{n, \bar{\pi}}(\mathcal{C}_N(\pi', \pi)), \quad (78)$$

and $\gamma = 2N^c$.

As an immediate consequence of Lemma H.3 and union bound, we have the following:

Lemma L.1. Fix $\delta \in (0, 1)$, and define $\varepsilon_{\text{stat}} = \frac{16 \log(4|\Pi|/\delta)}{n}$. With probability $1 - \delta$, the following holds simultaneously:

(1) For all $\pi, \pi' \in \Pi$, it holds that

$$\begin{aligned} 2\mathbb{P}_{\mu, \pi_D}(\mathcal{C}_N(\pi', \pi)) + \varepsilon_{\text{stat}} &\geq \mu_n(\mathcal{C}_N(\pi', \pi)) \geq \frac{1}{2}\mathbb{P}_{\mu, \pi_D}(\mathcal{C}_N(\pi', \pi)) - \varepsilon_{\text{stat}}, \\ 2\mathbb{P}_{n, \pi_D}(\mathcal{C}_N(\pi', \pi)) + \varepsilon_{\text{stat}} &\geq \mu_n(\mathcal{C}_N(\pi', \pi)) \geq \frac{1}{2}\mathbb{P}_{n, \pi_D}(\mathcal{C}_N(\pi', \pi)) - \varepsilon_{\text{stat}}. \end{aligned}$$

(2) For any $\pi \in \Pi$, it holds that $\mathcal{E}_{\gamma, \mu_n}(\pi_D \parallel \pi) \leq 2\mathcal{E}_{\gamma, \mu}(\pi_D \parallel \pi) + \varepsilon_{\text{stat}}$.

In the following, we fix $\delta \in (0, 1)$ and condition on the success event of Lemma L.1. Let $\bar{\pi} \in \Pi$ denote some policy for which $\varepsilon_{\text{apx}} = \mathcal{E}_{\gamma, \mu}(\pi_D \parallel \bar{\pi})$. We denote $\varepsilon'_{\text{apx}} = \mathcal{E}_{\gamma, \mu_n}(\pi_D \parallel \bar{\pi})$, and by Lemma L.1, we have $\varepsilon'_{\text{apx}} \leq 2\varepsilon_{\text{apx}} + \varepsilon_{\text{stat}}$.

Then, for any $\pi' \in \Pi$,

$$\begin{aligned} \mathcal{L}(\bar{\pi}, \pi') &\leq 2\mathbb{P}_{n, \pi_D}(\mathcal{C}_N(\pi', \bar{\pi})) - 2\gamma\mathbb{P}_{n, \bar{\pi}}(\mathcal{C}_N(\pi', \bar{\pi})) + \varepsilon_{\text{stat}} \\ &\leq 2\mathcal{E}_{\gamma, \mu_n}(\pi_D \parallel \bar{\pi}) + \varepsilon_{\text{stat}} = \varepsilon'_{\text{apx}} + \varepsilon_{\text{stat}}. \end{aligned}$$

where the first inequality uses Lemma L.1, and the second inequality uses Eq. (72).

Therefore, we have

$$\max_{\pi' \in \Pi} \mathcal{L}(\hat{\pi}, \pi') = \min_{\pi \in \Pi} \max_{\pi' \in \Pi} \mathcal{L}(\pi, \pi') \leq \max_{\pi' \in \Pi} \mathcal{L}(\bar{\pi}, \pi') \leq \varepsilon_{\text{stat}} + \varepsilon'_{\text{apx}}.$$

In particular, we know $\mathcal{L}(\hat{\pi}, \bar{\pi}) \leq \varepsilon_{\text{stat}} + \varepsilon'_{\text{apx}}$. Then, we can bound

$$\begin{aligned} \mu_n(\mathcal{C}_N(\bar{\pi}, \hat{\pi})) - \mathcal{L}(\hat{\pi}, \bar{\pi}) &= 2\gamma\mathbb{P}_{n, \hat{\pi}}(\mathcal{C}_N(\bar{\pi}, \hat{\pi})) \\ &\leq \frac{2\gamma}{N}\mathbb{P}_{n, \bar{\pi}}(\mathcal{C}_N(\bar{\pi}, \hat{\pi})) \\ &\leq \frac{2\gamma}{N} \left[\frac{\gamma}{2}\mathbb{P}_{n, \pi_D}(\mathcal{C}_N(\bar{\pi}, \hat{\pi})) + \varepsilon'_{\text{apx}} \right] \\ &\leq \frac{2\gamma}{N} [\gamma(\mu_n(\mathcal{C}_N(\bar{\pi}, \hat{\pi})) + \varepsilon_{\text{stat}}) + \varepsilon'_{\text{apx}}], \end{aligned}$$

where the second inequality uses Eq. (73): $\mathbb{P}_{n, \bar{\pi}}(E) - \frac{\gamma}{2}\mathbb{P}_{n, \pi_D}(E) \leq \mathcal{E}_{M, \mu_n}(\pi_D \parallel \bar{\pi}) = \varepsilon'_{\text{apx}}$ for any event E , and the third inequality uses Lemma L.1. Therefore, using $N \geq 4\gamma^2$, we know $\mu_n(\mathcal{C}_N(\bar{\pi}, \hat{\pi})) \leq 5\varepsilon_{\text{stat}} + 2\varepsilon'_{\text{apx}}$. Using Lemma L.1 again, we have

$$\text{Cov}_N^{\pi_D}(\bar{\pi} \parallel \hat{\pi}) = \mathbb{P}_{\mu, \pi_D}(\mathcal{C}_N(\bar{\pi}, \hat{\pi})) \leq 2\mu_n(\mathcal{C}_N(\bar{\pi}, \hat{\pi})) + 2\varepsilon_{\text{stat}} \leq 12\varepsilon_{\text{stat}} + 4\varepsilon'_{\text{apx}}.$$

By [Eq. \(71\)](#), it holds that

$$\text{Cov}_{2N\gamma}(\pi) = \text{Cov}_{2N\gamma}^{\pi_D}(\pi_D \parallel \pi) \leq \text{Cov}_{2\gamma}^{\pi_D}(\pi_D \parallel \bar{\pi}) + \text{Cov}_N^{\pi_D}(\bar{\pi} \parallel \pi),$$

and we also have $\text{Cov}_{2\gamma}^{\pi_D}(\pi_D \parallel \bar{\pi}) \leq \varepsilon_{\text{apx}}$ by [Eq. \(74\)](#). Combining the inequalities above, we can conclude that

$$\text{Cov}_{2N\gamma}(\hat{\pi}) \leq \text{Cov}_N^{\pi_D}(\bar{\pi} \parallel \pi) + \varepsilon_{\text{apx}} \leq 12\varepsilon_{\text{stat}} + 4\varepsilon'_{\text{apx}} + \varepsilon_{\text{apx}}.$$

Finally, using [Lemma H.3](#), we have $\varepsilon'_{\text{apx}} \leq 2\varepsilon_{\text{apx}} + \varepsilon_{\text{stat}}$. This is the desired upper bound. $\square$