

Soft Separation for Adaptive Robust Optimization

Qingyuan Xu, Ruiwei Jiang

Department of Industrial and Operations Engineering, University of Michigan, {qyxu, ruiwei}@umich.edu

We propose an algorithmic framework for solving adaptive robust optimization with provable guarantees on both tractability and solution accuracy. The framework introduces soft separation, a probabilistic mechanism for identifying worst-case uncertainty realizations via a time-inhomogeneous Markov chain. Rather than solving an exact separation problem in each iteration, which is intractable in general, the chain carries adversarial information across iterations, co-evolves with the optimization iterates, and recovers exact separation at terminal iterates with high probability. For continuous first-stage (here-and-now) decisions, we design a first-order method that uses soft separation to produce adaptive gradient estimates. Notably, we prove *polynomial-time convergence* in expectation to the global optimum. For mixed-integer here-and-now decisions, we embed soft separation within a branch-and-cut framework to generate valid cuts for the robust objective and obtain a high-probability certificate of global optimality. Numerical experiments demonstrate that the proposed methods scale favorably with problem dimension and scenario size relative to state-of-the-art approaches. The instances and code are available online at <https://github.com/xuqy2002/SoftSeparationARO>.

Key words: Robust optimization, soft separation, Markov chain

1. Introduction

Adaptive robust optimization (ARO) is an important framework for optimization under uncertainty. To hedge against adversarial uncertainty, it adapts to the uncertain realization through adjustable second-stage (also known as “wait-and-see”) recourse decisions, which allow ARO to provide reliable decisions while remaining less conservative than static robust optimization alternatives. ARO finds wide-ranging applications in, e.g., healthcare (Kong et al. 2013, Qi 2017, Ryu and Jiang 2025), power systems (Jiang et al. 2011, An and Zeng 2014, Moreira et al. 2024), and supply chains (Lu et al. 2015, Cheng et al. 2018, Qi et al. 2024). Formally, ARO is formulated as

$$F^* := \min_{x \in \mathcal{X}} F(x) := \left\{ f(x) + \max_{\xi \in \Xi} Q(x, \xi) \right\}, \quad (\text{ARO})$$

where f denotes a cost function, $x \in \mathcal{X} \subseteq \mathbb{R}^{d_x}$ denotes the first-stage (or “here-and-now”) decision, and $\xi \in \Xi \subseteq \mathbb{R}^{d_\xi}$ denotes the uncertain parameter with Ξ representing an uncertainty set. Both x and ξ may be continuous or discrete. For fixed x and ξ , the recourse function is defined through a convex conic program

$$Q(x, \xi) := \min_y b^\top y \quad \text{s.t.} \quad Gy + Ex + U\xi - h \in \mathcal{K}, \quad (1)$$

where $y \in \mathbb{R}^{d_y}$ denotes the wait-and-see recourse decisions and $\mathcal{K} \subseteq \mathbb{R}^m$ denotes a convex cone. For example, when $\mathcal{K} := \mathbb{R}_+^m$, the recourse formulation (1) reduces to a linear program and (ARO) is a two-stage robust *linear* program. Despite wide applications, (ARO) is computationally challenging (Minoux 2011), largely because of the inner max-min problem. For fixed x , this problem seeks to identify a worst-case scenario $\xi^* \in \arg \max_{\xi \in \Xi} Q(x, \xi)$, or equivalently, exactly separate the epigraph of the worst-case recourse function $\max_{\xi \in \Xi} Q(x, \xi)$. Exact solution algorithms for (ARO) solve the problem *repeatedly*. For example, the column-and-constraint generation (C&CG) algorithm (Zeng and Zhao 2013) repeatedly refines a relaxation to (ARO) by appending new ξ^* , together with the corresponding recourse variables and constraints. Likewise, the Benders decomposition algorithm (see, e.g., Jiang et al. 2011) repeatedly separates the worst-case recourse function by a supporting hyperplane of $Q(x, \xi^*)$, also known as an “optimality cut.”

Unfortunately, the exact separation problem is usually intractable. For example, when $\mathcal{K} = \mathbb{R}_+^m$ and the uncertainty set Ξ admits a parametric form such as a polyhedron or an ellipsoid, dualizing the recourse linear program (1) recasts the problem as a bilinear program, which is generally NP-complete (Bennett and Mangasarian 1993). In addition, when Ξ is a finite set of scenarios with no clear structure, the exact separation boils down to enumerating Ξ and solving the corresponding recourse problem (1) for each scenario, which quickly becomes prohibitive as the scenario size increases (see Section 5 for numerical demonstration).

We argue, however, that an exact separation is not necessary in every step of an iterative algorithm like C&CG or Benders. Early in the iterative algorithm, the incumbent solution can be far from optimum and an adversarial (but not necessarily the worst) scenario can already reveal useful directions for improvement. Moreover, adversarial scenarios often persist across nearby incumbent solutions. Hence, the adversarial scenario found for an incumbent solution need not be discarded immediately after the incumbent is updated. Instead, it is computationally much cheaper to carry forward and refine it for future incumbents.

These observations motivate **soft separation**. Instead of solving a fresh exact separation problem in each iteration, we track adversarial scenarios through a time-*inhomogeneous* Markov chain over the uncertainty set. At iteration n , the transition kernel uses the current incumbent x_n , while the chain state records an adversarial scenario. As x_n changes, the chain updates this state rather than restarting the adversarial search from scratch. Each separation therefore reduces to a simple scenario-wise update, while the Markov chain carries adversarial information across iterations. To back up this idea rigorously, we show that, when the x_n sequence changes in a controlled way and the soft separation is appropriately designed, the Markov chain identifies adversarial scenarios for all later-stage iterates of x_n with high probability. Then, we apply this soft-separation oracle in two settings: for x being continuous, the oracle produces Markovian gradient estimates for a robust

surrogate; and for x being mixed-integer, the oracle generates valid cuts for the worst-case recourse function in a branch-and-cut framework. Notably, for continuous x , we prove polynomial-time convergence of x_n , in expectation, to the global optimum of (ARO). To our best knowledge, this provides the first *tractable* and *tight* approximation algorithm for general (ARO).

1.1. Literature Review

Robust optimization (RO) models can be categorized by how the uncertainty set Ξ is characterized. The majority of the literature focuses on parametric uncertainty sets such as polyhedra (see, e.g., Bertsimas and Sim 2004) or ellipsoids (see, e.g., Ben-Tal and Nemirovski 1999). These uncertainty sets have good interpretation and provide robustness guarantees that, with high probability, any feasible solution to a RO model remains feasible when exposed to out-of-sample random realizations of ξ . Such guarantees usually rely on distributional assumptions of ξ such as independence (among all ξ -entries), light-tailedness, symmetry, or boundedness (see, e.g., Table 4 of Bertsimas et al. 2021, for a summary). In addition, Ξ can be characterized by independent scenarios of ξ drawn from its (possibly latent) probability distribution. This gives rise to a finite Ξ and the ensuing (ARO) admits both a priori (Campi and Garatti 2008) and a posteriori (Campi and Garatti 2018) robustness guarantees. Different from the parametric uncertainty sets, the robustness guarantees herein are free of distributional assumptions, except for the independence of the scenarios in Ξ . This paper seeks to develop algorithms for both parametric and scenario uncertainty sets.

Existing algorithmic work for solving (ARO) falls into two streams, with the first aiming to solve (ARO) to global optimum and solve the exact separation more efficiently, and the second resorting to tractable approximations of (ARO). When Ξ admits a parametric form such as a polyhedron or an ellipsoid, the C&CG algorithm by Zeng and Zhao (2013) solves a relaxation of (ARO) by considering a subset of ξ -scenarios. To achieve global optimum, C&CG iteratively solves the exact separation problem and appends the ensuing ξ^* to this subset. Despite the challenges of the exact separation, C&CG has demonstrated empirical success in real-world engineering systems such as power grids (see, e.g., Yuan et al. 2016), transportation (Huang et al. 2023), and healthcare (Ji et al. 2022). Recently, Tsang et al. (2023) proposed a variant of C&CG that solves the (ARO) relaxations inexactly while still maintaining a finite convergence to the global optimum. In addition, variants of C&CG have been applied to solve (ARO) with mixed-integer recourse (Zhao and Zeng 2012) and decision-dependent uncertainty (Zeng and Wang 2022). As compared to these exact algorithms, this paper studies solution approaches with tractability guarantee, e.g., converging to ϵ -optimality of (ARO) within a number of iterations that scales polynomially in $1/\epsilon$.

On the other hand, to address the computational challenge of (ARO), decision rules have been proposed to restrict the recourse variables y in (1) as functions of ξ . For example, Bertsimas and

Goyal (2013) study a static policy, in which y is restricted to be a constant function of ξ . This casts (ARO) as a one-stage robust optimization problem and regains computational tractability (see, e.g., Ben-Tal and Nemirovski 1999, Bertsimas and Sim 2004). Later, Han et al. (2023) extend the static policy to a piecewise static one and retain good interpretability. Moreover, Ben-Tal et al. (2004) propose an affine decision rule (ADR), which restricts y to be an affine function of ξ , and subsequent studies generalize ADR through lifting (see, e.g., Bertsimas et al. 2019), Fourier-Motzkin elimination (Zhen et al. 2018), and to piecewise affine variants (see, e.g., Chen et al. 2008, Georghiou et al. 2020). As in static policies, ADR and its generalizations produce conservative but tractable approximations for (ARO). These policies generally do not solve (ARO) to global optimum, with a few exceptions, e.g., when Ξ is a standard simplex (Bertsimas and Goyal 2012) or a box (Zhen et al. 2018) or when (ARO) fulfills additional structural and technical assumptions (Georghiou et al. 2026). Nevertheless, these policies admit approximation guarantees when, e.g., all entries of ξ are nonnegative (Bertsimas et al. 2011, Bertsimas and Bidkhori 2015) or all recourse matrix entries are nonnegative (El Housni and Goyal 2021, El Housni et al. 2024). As compared to these decision rules, our algorithms seek to find near-optimal solutions to a general (ARO) model.

We summarize our main contributions as follows.

1. For (ARO) with continuous decision variables and a scenario uncertainty set Ξ , we adopt a first-order algorithm that approximates the gradients of the robust objective using a time-inhomogeneous Markov chain over Ξ , thereby avoiding exhaustive exact separation at each iteration. We prove polynomial-time convergence, in expectation, to the global optimum of (ARO).
2. We formalize this mechanism as a soft separation framework and establish probabilistic guarantees for identifying adversarial scenarios. We then extend the framework beyond finite scenarios to general uncertainty sets.
3. For (ARO) with mixed-integer decisions, we embed soft separation within a branch-and-cut algorithm to produce deterministically valid cuts for the robust objective. Leveraging the probabilistic guarantees of soft separation, we obtain a high-probability certificate of global optimality.
4. We demonstrate the proposed algorithms in extensive numerical experiments spanning continuous decisions and mixed-integer decisions, as well as scenario, ellipsoidal, and budget uncertainty sets. Across these settings, the proposed methods provide near-optimal solutions to (ARO) and exhibit favorable scalability.

The remainder of the paper is organized as follows. Section 2 studies the first-order algorithm to solve (ARO) with continuous decision variables and a scenario uncertainty set. The proposed gradient estimation method is formalized as the soft separation framework in Section 3 and then applied to (ARO) with mixed-integer decision variables in Section 4. We report the numerical experiment results in Section 5. We present all proofs in Appendix EC.4.

Notation: we denote by $\mathbf{0}, \mathbf{1}$ all-zero and all-one vectors, by Id an identity matrix, by $\mathbb{N}_+$ the nonnegative integers, by $[n] := \{1, \dots, n\}$ the set of running indices for $n \in \mathbb{N}_+$, by $\mathbb{B}^d$ the unit ℓ_2 -norm ball in $\mathbb{R}^d$, by $\mathbb{S}^{d-1}$ the sphere of $\mathbb{B}^d$, and by $O(\cdot)$ the big-O notation with $f(n) = O(g(n))$ indicating that there exists an $M > 0$ such that $|f(n)| \leq Mg(n)$ for all sufficiently large n .

2. Continuous Decisions and Scenario Uncertainty

We start with a baseline (ARO) model with continuous decision variables x and a scenario uncertainty set $\Xi := \{\xi^{(i)} : i \in \mathcal{I}\}$, where $\mathcal{I}$ is an index set. This setup is without loss of generality (WLOG) when Ξ is characterized by scenarios of ξ (Campi and Garatti 2008, 2018), or when Ξ is polyhedral because $Q(x, \xi)$ is convex in ξ , implying that Ξ can be replaced by the set of its extreme points WLOG. For ease of exposition, we make the following assumptions throughout this section.

ASSUMPTION 1 (Continuous x and cost function). *The set $\mathcal{X} \subseteq \mathbb{R}^{d_x}$ is compact and convex with $\|x - x'\|_1 \leq D$ for all $x, x' \in \mathcal{X}$. The first-stage cost function $f : \mathcal{X} \rightarrow \mathbb{R}$ is convex and L_f -Lipschitz on an open neighborhood of $\mathcal{X}$.*

ASSUMPTION 2 (Relatively complete recourse). *For every $(x, \xi) \in \mathcal{X} \times \Xi$, the recourse function $Q(x, \xi)$ is finite, that is, $|Q(x, \xi)| < \infty$.*

ASSUMPTION 3 (Linear recourse problem). $\mathcal{K} = \mathbb{R}_+^m$.

Assumptions 1–2 are standard in the literature of ARO. In particular, when (ARO) is a two-stage robust linear program with $f(x) := c^\top x$, Assumption 1 holds trivially with $L_f = \|c\|$. We make Assumption 3 for ease of presenting the main ideas without technical clutter, and we extend the results of this section to the general conic case in Appendix EC.1. In general, $Q(x, \xi)$ lacks smoothness in x (and differentiability, even under Assumption 3). We therefore work with its subgradients obtained from dual optimal solutions of the recourse problem. As a preliminary, the following lemma recalls relevant properties of $Q(x, \xi)$.

LEMMA 1. *For each scenario $\xi^{(i)}$, the function $Q(x, \xi^{(i)})$ admits a representation*

$$Q(x, \xi^{(i)}) = \max_{\substack{\pi \geq 0 \\ G^\top \pi = b}} \left\{ \pi^\top (h - Ex - U \xi^{(i)}) \right\}. \quad (2)$$

Let $\pi^*(x, i)$ denote a maximizer of (2), chosen to minimize $\|E^\top \pi\|_1$ in case of alternative dual optimal solutions. Then, it holds that $-E^\top \pi^*(x, i) \in \partial_x Q(x, \xi^{(i)})$. Furthermore, there exists a constant $L_Q > 0$ such that, with $L'_Q := L_Q \|E\|$ and $L'_\xi := L_Q \|U\|$, the following hold:

- (1) Lipschitz continuity in x : For every $\xi \in \Xi$, $Q(\cdot, \xi)$ is L'_Q -Lipschitz continuous on $\mathcal{X}$;
- (2) Lipschitz continuity in ξ : For every $x \in \mathcal{X}$, $Q(x, \cdot)$ is L'_ξ -Lipschitz continuous on Ξ ;
- (3) Bounded subgradient selection: For every $x \in \mathcal{X}$ and every scenario $\xi^{(i)}$, the selected subgradient is measurable and satisfies $\|E^\top \pi^*(x, i)\| \leq L'_Q$.

We note that $\pi^*(x, i)$ in Lemma 1 can be found by solving an auxiliary linear program, which minimizes $\|E^\top \pi\|_1$ over the polyhedron $\{\pi \geq 0 : G^\top \pi = b, \pi^\top (h - Ex - U \xi^{(i)}) = Q(x, \xi^{(i)})\}$, where $Q(x, \xi^{(i)})$ has been pre-computed in the linear program (2).

2.1. Mellowmax Approximation and Algorithm

Exact separation is challenging because evaluating $\max_{\xi \in \Xi} Q(x, \xi)$ demands maximizing a convex function. To this end, we propose the following Mellowmax approximation to “soften” the maximization operator and connect it with an expectation operator. For $w > 0$, we define

$$\mathcal{M}_w(x) := w \log \left(\frac{1}{|\mathcal{I}|} \sum_{i=1}^{|\mathcal{I}|} \exp \left\{ \frac{Q(x, \xi^{(i)})}{w} \right\} \right).$$

By construction, $\mathcal{M}_w$ replaces the maximization operator in the exact separation with a log-sum-exponential form. The following lemma quantifies its approximation error.

LEMMA 2 (Uniform error bound of Mellowmax, Asadi and Littman (2017)). *For every $x \in \mathcal{X}$ and $w > 0$, it holds that*

$$\max_{\xi \in \Xi} Q(x, \xi) - w \log |\mathcal{I}| \leq \mathcal{M}_w(x) \leq \max_{\xi \in \Xi} Q(x, \xi).$$

Moreover, $\lim_{w \rightarrow 0^+} \mathcal{M}_w(x) = \max_{\xi \in \Xi} Q(x, \xi)$.

Lemma 2 suggests that $\mathcal{M}_w(x)$ approximates $\max_{\xi \in \Xi} Q(x, \xi)$ to any prescribed accuracy by choosing a sufficiently small w . Notably, the approximation error depends *logarithmically* on $|\mathcal{I}|$, making $\mathcal{M}_w(x)$ appealing for large uncertainty sets. We therefore consider the following optimization problem,

$$\min_{x \in \mathcal{X}} F_w(x) := f(x) + w \log \left(\frac{1}{|\mathcal{I}|} \sum_{i=1}^{|\mathcal{I}|} \exp \left\{ \frac{Q(x, \xi^{(i)})}{w} \right\} \right) \quad (\text{Approx})$$

using $F_w(x)$ as a surrogate of the original robust objective function in (ARO). Notice that F_w is convex but need not be differentiable on $\mathcal{X}$. Hence, a natural thought is to apply the standard projected subgradient method to (Approx) and generate a near-optimal solution to (ARO). However, an examination of a subgradient reveals that

$$\begin{aligned} \partial F_w(x) &\ni z_w(x) := g_f(x) + \sum_{i \in \mathcal{I}} \frac{\exp(Q(x, \xi^{(i)})/w)}{\sum_{j \in \mathcal{I}} \exp(Q(x, \xi^{(j)})/w)} (-E^\top \pi^*(x, i)) \\ &= \sum_{i \in \mathcal{I}} (g_f(x) - E^\top \pi^*(x, i)) p_x(i), \end{aligned}$$

where $g_f(x) \in \partial f(x)$ is a bounded measurable subgradient selection and

$$p_x(i) := \frac{\exp(Q(x, \xi^{(i)})/w)}{\sum_{j \in \mathcal{I}} \exp(Q(x, \xi^{(j)})/w)}.$$

A direct implementation would necessitate enumerating all scenarios in Ξ at every step of the subgradient method, which is prohibitively expensive. If we resort to a stochastic approximation, then in each iteration n we sample an $I_{n+1} \in \mathcal{I}$ conditionally on $\mathcal{F}_n$ according to $\mathbb{P}(I_{n+1} = i \mid \mathcal{F}_n) = p_{x_n}(i)$ and generate a random proxy $Z_{n+1} := g_f(x_n) - E^\top \pi^*(x_n, I_{n+1})$ for the selected subgradient. Here, x_n is the decision incumbent at iteration n and $\mathcal{F}_n$ denotes the natural filtration generated by the algorithm. By construction, $\mathbb{E}[Z_{n+1} \mid \mathcal{F}_n] \in \partial F_w(x_n)$, so that $x_{n+1} = \Pi_{\mathcal{X}}(x_n - \alpha_n Z_{n+1})$ is a stochastic projected subgradient update with an unbiased subgradient estimator.

Unfortunately, this construction does not resolve the computational challenge for two reasons. First, a direct sampling from p_{x_n} requires evaluating all scenario weights and consequently computing all $Q(x, \xi^{(i)})$. Second, the target distribution p_{x_n} changes with the iterate x_n . Thus, even when a suitable sampling method is available, we may need to rerun it in every iteration and repeating this procedure offsets much of the computational gain.

As an alternative, we propose to not insist on an exact and unbiased gradient estimate in every iteration. Instead, we maintain an inexact estimate of $p_{x_n}(i)$ using a time-inhomogeneous Markov chain. This chain favors scenarios $i \in \mathcal{I}$ with larger $Q(x_n, \xi^{(i)})$ values, which imitate $p_{x_n}(i)$. In addition, it co-evolves with x_n throughout the algorithm and coincides with $p_{x_n}(i)$ in the limit. Specifically, we adopt a sample-rejection procedure with a proposal distribution $q(j|i)$ and an acceptance distribution $A_x(i, j)$. Iteratively, when the Markov chain is at state $i \in \mathcal{I}$, we first sample a proposal state $j \in \mathcal{I}$ according to $q(\cdot|i)$ and then accept or reject the new state j with probability $A_x(i, j)$. This gives rise to a parameterized Markov kernel $\{P_x : x \in \mathcal{X}\}$ with, for all $i, j \in \mathcal{I}$,

$$P_x(i, j) := \begin{cases} q(j|i) A_x(i, j), & \text{if } j \neq i \\ 1 - \sum_{k \neq i} q(k|i) A_x(i, k), & \text{if } j = i. \end{cases}$$

For fixed x , the Markov chain induced by P_x is irreducible whenever $q(j|i) > 0$ for all $i, j \in \mathcal{I}$ and aperiodic if $q(i|i) > 0$. In particular, for a finite $\mathcal{I}$, irreducibility guarantees the existence of a unique stationary distribution, and we shall henceforth assume that such uniqueness holds. We give two examples for such a Markov chain.

EXAMPLE 1 (METROPOLIS–HASTINGS (MH) ACCEPTANCE). For fixed proposal distribution $q(\cdot|\cdot)$ and incumbent x , the MH acceptance probability puts

$$A_x^{\text{MH}}(i, j) := \min \left\{ 1, \frac{p_x(j) q(i|j)}{p_x(i) q(j|i)} \right\} = \min \left\{ 1, \exp \left(\frac{Q(x, \xi^{(j)}) - Q(x, \xi^{(i)})}{w} \right) \cdot \frac{q(i|j)}{q(j|i)} \right\}, \quad \forall i, j \in \mathcal{I}.$$

It follows that $p_x(i)P_x(i, j) = p_x(j)P_x(j, i)$ for all i, j and hence p_x is the stationary distribution of the Markov chain. $\square$

EXAMPLE 2 (BARKER ACCEPTANCE). A smoother acceptance rule proposed by Barker puts

$$A_x^B(i, j) := \frac{r_x(i, j)}{1 + r_x(i, j)}, \quad \text{where } r_x(i, j) := \frac{p_x(j) q(i | j)}{p_x(i) q(j | i)} = \exp\left(\frac{Q(x, \xi^{(j)}) - Q(x, \xi^{(i)})}{w}\right) \cdot \frac{q(i | j)}{q(j | i)}$$

for all $i, j \in \mathcal{I}$. It follows that $p_x(i)P_x(i, j) = p_x(j)P_x(j, i)$ and so the Markov chain admits p_x as the stationary distribution. $\square$

Examples 1–2 demonstrate that the proposed procedure does not demand evaluating all the recourse functions $Q(x, \xi^{(i)})$ in each iteration. In fact, under either MH or Barker acceptance, one evaluates the function exactly twice, one for the current state i and the other for the proposed state j . This equals solving two linear programs (2) and hence is computationally much cheaper than the exact gradient estimation or its stochastic approximation. We shall rely on either MH or Barker acceptance distributions in the remainder of the paper.

In addition, the proposed procedure is different from a conventional Markov chain Monte Carlo procedure, which would freeze x_n and apply P_{x_n} for many transitions, possibly after a burn-in period, to obtain a sample close to following the stationary distribution p_{x_n} . Since p_{x_n} changes with x_n , this inner sampling procedure would incur higher computational burden in every optimization iteration. Instead, we carry x_n and P_{x_n} forward and apply only a small number of transitions, possibly just one, before updating x_n . The ensuing Markov chain and the optimization iterates therefore evolve on the same time scale. Concretely, after sampling a scenario I_{n+1} according to P_{x_n} , we form

$$Z_{x_n}(I_{n+1}) := g_f(x_n) - E^\top \pi^*(x_n, I_{n+1})$$

and update

$$x_{n+1} = \Pi_{\mathcal{X}}(x_n - \alpha_{n+1} Z_{x_n}(I_{n+1})).$$

based on a certain step size α_{n+1} to be specified later. This yields a coupled framework in which the state is carried forward to identify an informative scenario, while the ensuing scenario-wise gradient updates the decision incumbent. We summarize the complete framework in Algorithm 1.

2.2. Theoretical Guarantee

Although the one-step Markov sample Z_{n+1} generally does not yield a conditionally unbiased subgradient estimate, our analysis shows that the resulting bias remains controllable because the chain tracks the evolving target distribution p_{x_n} sufficiently closely. Problems of this type, in which first-order estimates are constructed from observations generated by a decision-dependent Markov chain, arise naturally in policy-gradient reinforcement learning (Sutton and Barto 2018, Che et al. 2026), state-dependent performative prediction (Li and Wai 2022), and queueing and inventory systems (Li et al. 2026). Our problem of interest (ARO) does not by itself possess such an adaptive

Algorithm 1 Projected Stochastic Subgradient Descent with Soft Separation**Input:** Proposal kernel $q(\cdot | \cdot) > 0$, x_0 , parameters $w > 0$, iteration budget N , state space $\mathcal{I}$ **Output:** Weighted average iterate $\bar{x}_N$

- 1: $x \leftarrow x_0$, $I_0 \sim \text{Unif}\{\mathcal{I}\}$, $S \leftarrow 0$, $\bar{x} \leftarrow 0$
- 2: **for** $n = 0, \dots, N - 1$ **do**
- 3: $\alpha_{n+1} \leftarrow (n + 1)^{-1/2}$ and sample $J_n \sim q(\cdot | I_n)$
- 4: Evaluate $Q(x, \xi^{(I_n)})$ and $Q(x, \xi^{(J_n)})$ by solving (2)
- 5: Compute the MH or Barker acceptance probability $A_x(I_n, J_n)$
- 6: With probability $A_x(I_n, J_n)$, set $I_{n+1} \leftarrow J_n$
- 7: Obtain the optimal dual solution $\pi^*(x, \xi^{(I_{n+1})})$ from (2)
- 8: $Z_{x_n}(I_{n+1}) \leftarrow g_f(x) - E^\top \pi^*(x, \xi^{(I_{n+1})})$
- 9: $x \leftarrow \Pi_{\mathcal{X}}(x - \alpha_{n+1} Z_{x_n}(I_{n+1}))$
- 10: $\begin{cases} S \leftarrow S + \alpha_{n+1}, \bar{x} \leftarrow \bar{x} + \alpha_{n+1}x & \text{(Step-size Average)} \\ S \leftarrow S + (n + 1), \bar{x} \leftarrow \bar{x} + (n + 1)x & \text{(Linear Average)} \end{cases}$
- 11: **return** $\bar{x}_N \leftarrow \bar{x}/S$

data structure, nor does it admit a direct stochastic subgradient implementation of this form. Rather, this structure emerges from the specific soft-separation mechanism developed in the last subsection.

We establish theoretical guarantee in two major steps. We first use a Poisson-equation decomposition to control the cumulative error induced by the time-inhomogeneous Markov chain. Then, we combine this error bound with the first-order recursion to establish convergence of Algorithm 1.

To proceed, we first introduce the following lemma.

LEMMA 3 (Bounded subgradients of F_w). *Under Assumptions 1–2, F_w is convex on $\mathcal{X}$ and admits the subgradient selection $z_w(x) := \sum_{i \in \mathcal{I}} p_x(i) (g_f(x) - E^\top \pi^*(x, i)) \in \partial F_w(x)$. It holds that*

$$\|Z_x(i)\| \leq L_f + L'_Q, \quad \|z_w(x)\| \leq L_f + L'_Q.$$

Next, we analyze the error induced by the one-step Markov sample through a centered function-value difference. Fix an $x^* \in \arg \min_{x \in \mathcal{X}} F_w(x)$ and denote

$$e_{n+1} := Q(x_n, \xi^{(I_{n+1})}) - Q(x^*, \xi^{(I_{n+1})}) - \sum_{j \in \mathcal{I}} p_{x_n}(j) (Q(x_n, \xi^{(j)}) - Q(x^*, \xi^{(j)})),$$

the centered recourse-value difference. We recall that e_{n+1} is generally neither independent nor conditionally mean-zero unless the sampling distribution coincides with the stationary distribution

p_{x_n} . We analyze the transient behavior of e_{n+1} through the so-called fundamental matrix $(\text{Id} - P_x + \mathbf{1} p_x^\top)^{-1}$ of a Markov chain. For a finite and irreducible Markov kernel P_x , it holds that

$$(\text{Id} - P_x + \mathbf{1} p_x^\top)^{-1} - \mathbf{1} p_x^\top = \sum_{n=0}^{\infty} (P_x^n - \mathbf{1} p_x^\top),$$

where $P_x^n - \mathbf{1} p_x^\top$ evaluates the difference between the n -step Markov kernel and the stationary distribution and hence the right-hand side measures the accumulated (transient) deviation of the Markov chain from its stationarity. Hence, a bounded accumulated deviation demands us to control the magnifying power of the fundamental matrix, which we formalize below as an assumption.

ASSUMPTION 4 (Uniform norm bound of the fundamental matrix). *With respect to any fixed operator norm, e.g., $\|\cdot\|_{\infty \rightarrow \infty}$, it holds that*

$$\sup_{x \in \mathcal{X}} \|(\text{Id} - P_x + \mathbf{1} p_x^\top)^{-1}\| \leq C_*$$

for a constant $0 < C_* < \infty$.

The validity of Assumption 4 and the uniform norm bound C_* depend on P_x and hence on the proposal kernel $q(\cdot|\cdot)$. In other words, different proposal and acceptance distributions give rise to different C_* . Next, we recall a general approach to bounding C_* and then elaborate its applications to various proposal kernels.

LEMMA 4 (Doebelin minorization). *If P_x satisfies the Doebelin minorization condition*

$$P_x(i, j) \geq \gamma_x \nu_x(j) \quad \forall i, j \in \mathcal{I}$$

for some constant $\gamma_x > 0$ and probability measure ν_x on $\mathcal{I}$, then it holds that $\|(\text{Id} - P_x + \mathbf{1} p_x^\top)^{-1}\|_{\infty \rightarrow \infty} \leq 2/\gamma_x$.

EXAMPLE 3 (INDEPENDENT PROPOSAL). Consider an independent proposal with $q(j|i) \equiv q(j)$ and suppose that there exists a $c \geq 1$, independent of x , such that

$$c^{-1} p_x(j) \leq q(j) \leq c p_x(j), \quad \forall x \in \mathcal{X}, \forall j \in \mathcal{I}.$$

Then for all x and $i \neq j$, it holds that $r_x(i, j) = \frac{p_x(j)q(i)}{p_x(i)q(j)} \in [c^{-2}, c^2]$.

First, for the MH acceptance (see Example 1), we have $A_x^{\text{MH}}(i, j) = \min\{1, r_x(i, j)\} \geq \min\{1, c^{-2}\} = c^{-2}$. Then, $P_x(i, j) = q(j) A_x^{\text{MH}}(i, j) \geq (1/c^2) q(j)$ whenever $i \neq j$.

Second, for the Barker acceptance (see Example 2), we have $A_x^{\text{B}}(i, j) = \frac{r_x(i, j)}{1 + r_x(i, j)} \geq \frac{c^{-2}}{1 + c^{-2}} = \frac{1}{1 + c^2}$. Then, $P_x(i, j) = q(j) A_x^{\text{B}}(i, j) \geq (1/(1 + c^2)) q(j)$ whenever $i \neq j$.

Finally, for both MH and Barker acceptances, we have $0 \leq A_x^{\text{MH}}(i, j), A_x^{\text{B}}(i, j) \leq 1$ for all $i \neq j$. Hence, by definition of the transition kernel, $P_x(i, i) = 1 - \sum_{j \neq i} q(j) A_x(i, j) \geq 1 - \sum_{j \neq i} q(j) = q(i)$, where $A_x(i, j)$ refers to either $A_x^{\text{MH}}(i, j)$ or $A_x^{\text{B}}(i, j)$.

Therefore, the Doeblin minorization holds, uniformly on $\mathcal{X}$, for the MH acceptance with $\gamma_{\text{MH}} = 1/c^2$ and $\nu_{\text{MH}}(j) = q(j)$, and for the Barker acceptance with $\gamma_{\text{B}} = 1/(1+c^2)$ and $\nu_{\text{B}}(j) = q(j)$. It follows that $C_*^{\text{MH}} \leq 2c^2$ and $C_*^{\text{B}} \leq 2(1+c^2)$. $\square$

EXAMPLE 4 (APPROXIMATELY CORRECT PROPOSAL). Consider a proposal distribution that is approximately correct, in the sense that

$$\frac{1}{\kappa} \frac{p_x(i)}{p_x(j)} \leq \frac{q(i|j)}{q(j|i)} \leq \kappa \frac{p_x(i)}{p_x(j)} \quad \forall i, j \in \mathcal{I}$$

for a constant $\kappa \geq 1$. In other words, $q(\cdot | \cdot)$ is largely aligned with the target ratios. Then, $r_x(i, j) \in [1/\kappa, \kappa]$ and so $A_x^{\text{MH}}(i, j) \geq \frac{1}{\kappa}$, $A_x^{\text{B}}(i, j) \geq \frac{1}{1+\kappa}$. Define $q_{\min}(j) := \inf_{i \in \mathcal{I}} q(j | i) > 0$. Then, $P_x(i, i) \geq q(i | i)$ and whenever $i \neq j$, it holds that

$$P_x(i, j) \geq \kappa^{-1} q_{\min}(j) =: b_j^{\text{MH}} \quad (\text{MH})$$

$$P_x(i, j) \geq (1 + \kappa)^{-1} q_{\min}(j) =: b_j^{\text{B}}. \quad (\text{Barker})$$

Hence, the Doeblin minorization holds, uniformly on $\mathcal{X}$, for the MH acceptance with $\gamma_{\text{MH}} := \sum_j \min\{b_j^{\text{MH}}, q(j | j)\}$ and $\nu_{\text{MH}}(j) := \min\{b_j^{\text{MH}}, q(j | j)\}/\gamma_{\text{MH}}$, and for the Barker acceptance with $\gamma_{\text{B}} := \sum_j \min\{b_j^{\text{B}}, q(j | j)\}$ and $\nu_{\text{B}}(j) := \min\{b_j^{\text{B}}, q(j | j)\}/\gamma_{\text{B}}$. Therefore, $C_*^{\text{MH}} \leq 2/\gamma_{\text{MH}}$ and $C_*^{\text{B}} \leq 2/\gamma_{\text{B}}$. $\square$

EXAMPLE 5 (UNIFORM PROPOSAL). Consider a uniform proposal with $q(j | i) = |\mathcal{I}|^{-1}$ for all $i, j \in \mathcal{I}$. Denote $\Delta := \sup_{x \in \mathcal{X}} \max_{i, j \in \mathcal{I}} |Q(x, \xi^{(i)}) - Q(x, \xi^{(j)})|$. Then, $\Delta < \infty$ by Assumption 2 and $A_x^{\text{MH}}(i, j) \geq e^{-\Delta/w}$, $A_x^{\text{B}}(i, j) \geq (1 + e^{\Delta/w})^{-1}$. It follows that

$$P_x(i, j) \geq \frac{e^{-\Delta/w}}{|\mathcal{I}|} \quad (\text{MH}), \quad P_x(i, j) \geq \frac{(1 + e^{\Delta/w})^{-1}}{|\mathcal{I}|} \quad (\text{Barker}).$$

Therefore, the Doeblin minorization holds, uniformly on $\mathcal{X}$, for the MH acceptance with $\gamma_{\text{MH}} = e^{-\Delta/w}$ and $\nu_{\text{MH}}(j) = q(j)$, and for the Barker acceptance with $\gamma_{\text{B}} = (1 + e^{\Delta/w})^{-1}$ and $\nu_{\text{B}}(j) = q(j)$. It follows that $C_*^{\text{MH}} \leq 2e^{\Delta/w}$ and $C_*^{\text{B}} \leq 2(1 + e^{\Delta/w})$.

Furthermore, C_* can be bounded without using Δ . Define $p_{\max} := \sup_{x \in \mathcal{X}} \max_{i \in \mathcal{I}} p_x(i)$. Then, $A_x^{\text{MH}}(i, j) = \min\left\{1, \frac{p_x(j)}{p_x(i)}\right\} \geq p_x(j)/p_{\max}$ for all $i, j \in \mathcal{I}$. It follows that $P_x(i, j) \geq p_x(j)/(|\mathcal{I}|p_{\max})$, that is, the Doeblin minorization holds for the MH acceptance with $\gamma_{\text{MH}} = 1/(|\mathcal{I}|p_{\max})$ and $\nu_{\text{MH}}(j) = p_x(j)$. Likewise, it holds for the Barker acceptance with $\gamma_{\text{B}} = 1/(2|\mathcal{I}|p_{\max})$ and $\nu_{\text{B}}(j) = p_x(j)$. These strengthen the bounds for C_* to be $C_*^{\text{MH}} \leq \min\{2e^{\Delta/w}, 2|\mathcal{I}|p_{\max}\}$ and $C_*^{\text{B}} \leq \min\{2(1 + e^{\Delta/w}), 4|\mathcal{I}|p_{\max}\}$. $\square$

The independent proposal (Example 3) and the approximately correct proposal (Example 4) are useful when we can roughly locate the “bad” scenarios ξ with a large $Q(x, \xi)$, because their

ensuing $p_x(\cdot)$ values are significantly larger than those of the “good” scenarios. For example, if $Q(x, \xi)$ models the total cost of resource allocation (such as staffing and workforce scheduling) with random demands ξ , then $Q(x, \xi)$ is monotone (entrywise) in ξ and a scenario becomes worse as it increases (Ryu and Jiang 2025, Aykin 1996). In this case, a reasonable proposal distribution sets $q(j)$ (respectively, $q(j | \cdot)$) higher for a “bad” scenario $\xi^{(j)}$ that is entrywise maximal. This yields a small c (respectively, κ), which in turn implies a small C_* .

On the other hand, one can apply the uniform proposal (Example 5) without any prior knowledge about (ARO) or bounds learned from intermediate iterations of Algorithm 1. We adopt both uniform and approximately correct proposals for the numerical experiments reported in Section 5. In the following example, we discuss how to sample uniformly from Ξ (or its extreme points or boundary) for several commonly-applied uncertainty sets.

EXAMPLE 6 (UNIFORMLY SAMPLING AN UNCERTAINTY SET). 1. If Ξ is discrete and finite with $\Xi = \{\xi^{(i)} : i \in \mathcal{I}\}$, then sampling uniformly from Ξ is equivalent to doing so from $\mathcal{I}$.

2. Consider a box uncertainty set with $\Xi = \{\xi : \ell \leq \xi \leq u\}$ and we seek to sample uniformly from its extreme points. Notice that Ξ is scaled from an elementary ℓ_∞ -ball because $\Xi = \{\xi : \xi = \ell + Du, u \in [0, 1]^{d_\xi}\}$, where $D := \text{diag}(u - \ell)$. We first sample $u_i \sim \text{Bernoulli}(1/2)$ independently for each entry of u . Then, the ensuing $\xi = \ell + Du$ is a uniform sample from the Ξ -extreme points.

3. Consider a scaled ℓ_1 uncertainty set with $\Xi = \{\xi : \xi = c + Du, \|u\|_1 \leq 1\}$, where $D = \text{diag}(D_1, \dots, D_{d_\xi})$ and we seek to sample uniformly from its extreme points. To this end, we sample an $i \sim \text{Unif}\{1, \dots, d_\xi\}$ and the ensuing $\xi = c \pm De_i$ is a uniform sample from the Ξ -extreme points.

4. Consider a two-sided budget uncertainty set $\Xi = \{\xi : |\xi_i| \leq 1, \sum_{i=1}^{d_\xi} |\xi_i| \leq \Gamma\}$, where $\Gamma := k + \rho$ denotes the “budget of uncertainty” (Bertsimas and Sim 2004) with an integer $k \in [d_\xi]$ and a fractional $\rho \in [0, 1)$. In order to sample uniformly from the extreme points of Ξ , we first notice that every extreme point has exactly (i) k entries equal to $+1$ or -1 , (ii) an additional entry equal to $+\rho$ or $-\rho$, and (iii) all remaining entries equal to 0. Hence, we first sample a k -element subset $S \subseteq [d_\xi]$ uniformly, and then choose $j \in [d_\xi] \setminus S$ uniformly. Finally, we assign values to the $k + 1$ chosen entries uniformly from $\{+1, -1\}$ to obtain a random sub-vector $s \in \mathbb{R}^{k+1}$. Then, the ensuing ξ with

$$\xi_i = \begin{cases} s_i & \text{if } i \in S \\ \rho s_j & \text{if } i = j \\ 0 & \text{otherwise} \end{cases}$$

for all $i \in [d_\xi]$ is a uniform sample from the Ξ -extreme points.

5. Consider an ellipsoidal uncertainty set with $\Xi = \{\xi : \xi = c + H^{-1/2}u, \|u\|_2 \leq 1\}$ and we seek to sample uniformly from its boundary. We first generate a normalized $u = g/\|g\|_2$ from a Gaussian random vector $g \sim \mathcal{N}(0, I)$. Then, we accept u with probability $\sqrt{u^\top H u / \lambda_{\max}(H)}$, where $\lambda_{\max}(H)$

denotes the largest eigenvalue of the Hessian matrix H . We will either resample u in case it is rejected, or accept it and the ensuing $\xi = c + H^{-1/2}u$ is a uniform sample from the Ξ -boundary. $\square$

Having addressed the non-independent stochastic approximation error, we now establish the stability of the incumbent trajectory. To this end, we examine the cumulative descent error $\mathbb{E}[-2 \sum_{k=1}^n \alpha_k e_k]$. By the definition of subgradients and e_k , we have

$$\begin{aligned} Z_{x_{k-1}}(I_k)^\top (x_{k-1} - x^*) &\geq f(x_{k-1}) - f(x^*) + Q(x_{k-1}, \xi^{(I_k)}) - Q(x^*, \xi^{(I_k)}) \\ &= f(x_{k-1}) - f(x^*) + \sum_{i \in \mathcal{I}} p_{x_{k-1}}(i) (Q(x_{k-1}, \xi^{(i)}) - Q(x^*, \xi^{(i)})) + e_k \\ &\geq F_w(x_{k-1}) - F_w(x^*) + e_k, \end{aligned}$$

where $x^* \in \operatorname{argmin}_{x \in \mathcal{X}} F_w(x)$ and the last line follows from Jensen's inequality because the logarithm function is concave and $p_{x_{k-1}}(\cdot)$ is a probability distribution on $\mathcal{I}$:

$$\begin{aligned} F_w(x^*) - F_w(x_{k-1}) &= f(x^*) - f(x_{k-1}) + w \log \left(\sum_{i \in \mathcal{I}} p_{x_{k-1}}(i) \exp \left\{ \frac{Q(x^*, \xi^{(i)}) - Q(x_{k-1}, \xi^{(i)})}{w} \right\} \right) \\ &\geq f(x^*) - f(x_{k-1}) + \sum_{i \in \mathcal{I}} p_{x_{k-1}}(i) (Q(x^*, \xi^{(i)}) - Q(x_{k-1}, \xi^{(i)})). \end{aligned}$$

Thus, $-\alpha_k e_k$ enters the optimality gap $F_w(x_{k-1}) - F_w(x^*)$ through the Markov sample $Z_{x_{k-1}}(I_k)$, and a smaller cumulative error yields a tighter convergence bound. The next proposition bounds the cumulative descent error through the step sizes α_k .

PROPOSITION 1 (Cumulative descent error bound). *Suppose that the step sizes are non-increasing with $\alpha_k > 0$. Define constants $B_Z := 2C_* DL'_Q$, $L_Z := C_*(2L'_Q + 2D(L'_Q)^2/w) + 2C_*^2 DL'_Q (L_P + 2L'_Q/w)$, where L_P denotes the Lipschitz modulus of P_x such that $\sum_{j \in \mathcal{I}} |P_x(i, j) - P_{x'}(i, j)| \leq L_P \|x - x'\|$ for all $i \in \mathcal{I}$ (see Lemmas EC.3–EC.4). Then, for any $n \geq 1$, it holds that*

$$\mathbb{E} \left[-2 \sum_{k=1}^n \alpha_k e_k \right] \leq C_1 \sum_{k=1}^n \alpha_k^2 + C_0, \quad (3)$$

where nonnegative constants C_0, C_1 depend only on $(B_Z, L_Z, L_P, L_f, L'_Q)$ through

$$C_0 = 4B_Z \alpha_1, \quad C_1 = 2(L'_Q + L_f)(L_P B_Z + L_Z).$$

Proposition 1 implies that the cumulative descent error can be controlled by choosing the step sizes properly. Intuitively, the step size α_k bounds the change in the iterates $\|x_k - x_{k-1}\|$ and allows the Markov chain to carry the memory about adversarial scenarios forward. We are now ready to present the main convergence guarantee for Algorithm 1.

THEOREM 1. *Under Assumptions 1–4, suppose that the step sizes $\{\alpha_k\}$ are positive and nonincreasing. Then, for any $n \geq 1$, we have*

$$0 \leq \mathbb{E}[F_w(\bar{x}_n)] - F_w^* \leq \frac{\mathbb{E}[\|x_0 - x^*\|^2] + C_0 + (C_1 + 3(L_f + L'_Q)^2) \sum_{k=1}^n \alpha_k^2}{2 \sum_{k=1}^n \alpha_k},$$

where $F_w^* := \min_{x \in \mathcal{X}} F_w(x)$. Moreover, if we choose the step sizes as $\alpha_n \propto 1/\sqrt{n+1}$, then Algorithm 1 achieves a convergence rate to F_w^* in the order of $O((C_* + C_*^2/w)n^{-1/2} \log n)$ for the step-size average and $O((C_* + C_*^2/w)n^{-1/2})$ for the linear average, respectively.

Theorem 1 establishes polynomial convergence of Algorithm 1 to the global minimum of $F_w(x)$. But Lemma 2 implies that $F_w(\cdot)$ can approximate the true objective function of (ARO) to any prescribed accuracy by choosing w sufficiently small. Therefore, Algorithm 1 provides a near-optimal solution to (ARO) in polynomial time. Furthermore, we can iteratively adjust w in Algorithm 1, producing w_n in iteration n , to gradually obtain a tighter approximation $F_{w_n}(\cdot)$ for $F(\cdot)$ as $\bar{x}_n$ progresses. This adaptive variant of Algorithm 1 converges to F^* in polynomial time.

PROPOSITION 2. *Under Assumptions 1–4, define $I^*(x) := \arg \max_{i \in \mathcal{I}} Q(x, \xi^{(i)})$, $\tilde{\Delta}_n := \max_{I \in \mathcal{I}} Q(x_n, \xi^{(I)}) - \max_{j \notin I^*(x_n)} Q(x_n, \xi^{(j)})$, and suppose that there exists a $\Delta_* > 0$ such that $\Delta_* \leq \tilde{\Delta}_n$ for all sufficiently large n . At iteration $n \geq 0$, use step size $\alpha_n \propto n^{-1/2}$ and temperature parameter $w_n = (2\Delta_*)/(\log((n+2)|\mathcal{I}|^2))$ for the Mellowmax approximation in Algorithm 1. Then, the linear average $\bar{x}_n$ satisfies*

$$0 \leq \mathbb{E}[F(\bar{x}_n)] - F^* \leq O\left(\frac{1 + C_*^2 \log(n|\mathcal{I}|)}{\sqrt{n}}\right).$$

3. Soft Separation

The last section uses soft separation in one concrete role: generating first-order estimates for the Mellowmax surrogate. The soft-separation mechanism itself, however, is independent of the particular iterative algorithm. In this section, we abstract it from Algorithm 1 to establish method-independent guarantees for tracking a time-varying adversarial distribution and to identify adversarial scenarios for all later-stage incumbents. We then apply these results to generalize the convergence guarantee of Algorithm 1 beyond finite scenarios to general uncertainty sets.

3.1. The Soft-Separation Oracle

Recall that soft separation replaces the exact separation with a stateful process, which updates the state through a transition rule toward scenarios that are adversarial for the decision incumbent. To formalize this process, we first let $\mathcal{I}$ be a measurable state space equipped with its Borel σ -algebra $\mathcal{B}(\mathcal{I})$ and $I \mapsto \xi_I \in \Xi$ be a measurable mapping from states to uncertainty realizations. We

assume that the representation is sufficiently rich to preserve the worst-case recourse value, namely, $\sup_{I \in \mathcal{I}} Q(x, \xi_I) = \sup_{\xi \in \Xi} Q(x, \xi)$ for all $x \in \mathcal{X}$. Note here that $\mathcal{I}$ need not be finite or discrete as in Section 2. For example, $\mathcal{I}$ may be the boundary of a sphere.

EXAMPLE 7 (LINEAR RECOURSE WITH AN ELLIPSOIDAL UNCERTAINTY SET). Consider a linear recourse model with $\mathcal{K} = \mathbb{R}_+^m$, $Q(x, \xi) = \min_y \{b^\top y : Gy \geq h - Ex - U\xi\}$, and an ellipsoidal uncertainty set $\Xi = \{\bar{\xi} + Au : \|u\|_2 \leq 1\}$ with $A \in \mathbb{R}^{d_\xi \times d_u}$. Because $Q(x, \cdot)$ is convex, its maximum over Ξ is attained on the boundary of the ellipsoid. We may therefore parameterize the adversarial state by

$$\mathcal{I} := \mathbb{S}^{d_u-1}, \quad \xi_I := \bar{\xi} + AI. \quad \square$$

EXAMPLE 8 (QUADRATIC RECOURSE WITH AN ELLIPSOIDAL UNCERTAINTY SET). Consider a quadratic recourse model with $\mathcal{K} = \mathbb{R}_+^m$, $Q(x, \xi) = \min_y \{\frac{1}{2}y^\top Hy + (b + B\xi)^\top y : Gy \geq h - Ex - U\xi\}$, and an ellipsoidal uncertainty set $\Xi = \{\bar{\xi} + Au : \|u\|_2 \leq 1\}$, where $H \succeq 0$ and $A \in \mathbb{R}^{d_\xi \times d_u}$. Because $Q(x, \xi)$ is not necessarily convex in ξ , the worst-case scenario may appear anywhere in Ξ . Consequently, we may parameterize the adversarial state by

$$\mathcal{I} := \mathbb{B}^{d_u}, \quad \xi_I := \bar{\xi} + AI. \quad \square$$

Second, to rank the states at a given incumbent x , we assign each $I \in \mathcal{I}$ a measurable score $\hat{Q}_x : \mathcal{I} \rightarrow \mathbb{R}$. The score may be exact with $\hat{Q}_x(I) \equiv Q(x, \xi_I)$ or approximate with $\sup_{I \in \mathcal{I}} |Q(x, \xi_I) - \hat{Q}_x(I)| \leq \varepsilon$ for all $x \in \mathcal{X}$ and $\varepsilon \geq 0$ represents a uniform error bound. For example, the Moreau envelope $\hat{Q}_x(I) := \min_{v \in \mathbb{R}^{d_x}} \{Q(v, \xi_I) + (1/2\tau)\|v - x\|^2\}$ admits a uniform error bound $\varepsilon = (\tau/2)(L'_Q)^2$.

Soft separation uses $\hat{Q}_x$ to bias the adversarial state toward realizations with larger $Q(x, \xi)$ values. Specifically, let ν be a finite reference measure on $\mathcal{I}$ with full support (e.g., ν is the normalized Lebesgue measure on $\mathcal{I}$). For a temperature parameter $w > 0$, we approximate the worst-case recourse function $\max_{\xi \in \Xi} Q(x, \xi)$ by the Mellowmax function

$$\mathcal{M}_w^\nu(x) := w \log \int_{\mathcal{I}} \exp \left\{ \frac{\hat{Q}_x(I)}{w} \right\} \nu(dI).$$

$\mathcal{M}_w^\nu(\cdot)$ generalizes $\mathcal{M}_w(\cdot)$, which was defined in Section 2 for a finite $\mathcal{I}$, to a general $\mathcal{I}$. We show in EC.4.7 that $\mathcal{M}_w^\nu(\cdot)$ admits an analogous uniform error bound for approximating $\max_{\xi \in \Xi} Q(x, \xi)$ as in Lemma 2. Accordingly, the robust objective function of (ARO) is approximated by $F_w^\nu(x) := f(x) + \mathcal{M}_w^\nu(x)$. Under Assumption 5, if $\hat{Q}_x$ and f are differentiable then $\|\nabla_x \hat{Q}_x(I)\|$ is bounded uniformly and so $\nabla F_w^\nu(x) = \nabla f(x) + \int_{\mathcal{I}} \nabla_x \hat{Q}_x(I) p_x^w(dI)$, where the soft adversarial distribution p_x^w is defined through

$$\frac{dp_x^w}{d\nu}(I) = \frac{\exp(\hat{Q}_x(I)/w)}{\int_{\mathcal{I}} \exp(\hat{Q}_x(J)/w) d\nu(J)} \quad \forall I \in \mathcal{I}.$$

The distribution p_x^w assigns greater mass to states with larger scores, and a smaller value of w produces stronger concentration around the highest-scoring states. We remark that p_x^w is well-defined whenever ν is a finite measure and $\widehat{Q}_x(I)$ is bounded.

Unfortunately, p_x^w is challenging to sample from because of the integral $\int_{\mathcal{I}} \exp(\widehat{Q}_x(J)/w) d\nu(J)$. In addition, we need to adapt the sampling from p_x^w to the incumbent x , which may be updated frequently in decision-making. To address these, we pair p_x^w with

- a Markov transition kernel P_x on $\mathcal{I}$ such that p_x^w is an invariant distribution of P_x for any fixed $x \in \mathcal{X}$. For example, P_x may be constructed through a proposal-acceptance scheme as in Examples 1–5 or a Langevin diffusion-type dynamics (Roberts and Tweedie 1996). The proposal may also exploit the geometry of Ξ and information from the recourse function $Q(x, \xi)$ and the score $\widehat{Q}_x(I)$. When the adversarial search can be restricted to a smooth boundary, as in linear recourse with ellipsoidal uncertainty (Example 7), one may use geodesic or tangent-space proposals (Byrne and Girolami 2013, Zappa et al. 2018). For a full-dimensional convex uncertainty set (Example 8), convex-body random walks such as hit-and-run provide feasible proposals (Lovász and Vempala 2006). Primal-dual information from $Q(x, \xi)$ may further be used to construct gradient-informed proposals. For example, when a smooth approximation to $\widehat{Q}_x$ is available, one may employ Hamiltonian-type proposals (Neal 2011).

- an iterative optimization algorithm that generates a sequence $\{x_n : n \in \mathbb{N}_+\}$ of incumbents. For example, this algorithm may be the projected subgradient descent as in Algorithm 1 or a branch-and-cut algorithm (see Section 4).

We then formalize the soft-separation oracle as follows.

DEFINITION 1 (MARKOVIAN SOFT-SEPARATION ORACLE). Fix a temperature parameter $w > 0$, a score accuracy $\varepsilon \geq 0$, and a state representation $\xi_I : \mathcal{I} \rightarrow \Xi$. Given an incumbent sequence $\{x_n : n \in \mathbb{N}_+\}$, a *Markovian soft-separation oracle* is a time-inhomogeneous Markov process $\{I_n : n \in \mathbb{N}_+\}$ on $\mathcal{I}$ satisfying $I_{n+1} \sim P_{x_n}(I_n, \cdot)$, where, for every $x \in \mathcal{X}$, the transition kernel P_x admits the soft adversarial distribution p_x^w as an invariant distribution.

A key insight of the oracle is that it decouples adversarial relevance from exact identification in every iteration. Although the incumbent sequence and the adversarial process co-evolve, their interaction can be controlled: when the incumbent changes gradually, the oracle tracks the corresponding adversarial target and can approach, or recover with high probability, exact separation at later-stage iterates. We show these properties next.

3.2. Certification for Identifying Adversarial Scenarios

We denote by $(\mathcal{X}, \|\cdot\|)$ a metric space for incumbent solutions. Our analysis applies to both discrete (finite) and continuous (uncountable) state spaces. We start by imposing regularity conditions on the oracle.

ASSUMPTION 5 (Lipschitz score). *There exists an $L'_Q < \infty$ such that, for every $i \in \mathcal{I}$ and all $x, x' \in \mathcal{X}$,*

$$|\widehat{Q}_x(i) - \widehat{Q}_{x'}(i)| \leq L'_Q \|x - x'\|.$$

ASSUMPTION 6 (Lipschitz transition kernel). *There exists an L_P such that, for all $x, y \in \mathcal{X}$,*

$$\sup_{I \in \mathcal{I}} \|P_x(I, \cdot) - P_y(I, \cdot)\|_{\text{TV}} \leq L_P \|x - y\|,$$

where $\|P - Q\|_{\text{TV}} := \sup_{B \in \mathcal{B}(\mathcal{I})} |P(B) - Q(B)|$ denotes the total variation between measures P, Q .

ASSUMPTION 7 (Uniform mixing). *There exists an $\bar{\eta} \in [0, 1)$ such that, for every $x \in \mathcal{X}$, the Dobrushin coefficient of P_x satisfies*

$$\eta_{\text{Dob}}(P_x) := \sup_{I, J \in \mathcal{I}} \|P_x(I, \cdot) - P_x(J, \cdot)\|_{\text{TV}} \leq \bar{\eta}.$$

Assumption 5 holds for the recourse function $Q(x, \xi)$, i.e., for $\widehat{Q}_x(\cdot) \equiv Q(x, \cdot)$ (see Lemma 1, whose proof applies to a general $\mathcal{I}$) and consequently for its Moreau envelope as well. In addition, Assumption 6 follows naturally from Assumption 5. For example, for both MH and Barker acceptance, we show the Lipschitz continuity of P_x for a discrete and finite $\mathcal{I}$ in Lemmas EC.3–EC.4 and for a general $\mathcal{I}$ in Lemma EC.5 through Lipschitz scores. Finally, Assumption 7 enables uniform mixing of the Markov process characterized by P_x (see Dobrushin 1956) and it holds when P_x is induced by appropriate proposal and acceptance distributions. For example, $\eta_{\text{Dob}}(P_x)$ can be uniformly bounded when the proposal distribution is uniform on $\mathcal{I}$ and the acceptance distribution is either MH or Barker.

EXAMPLE 9 (UNIFORM MIXING WITH UNIFORM PROPOSAL). We continue Example 7 with $\mathcal{I} = \mathbb{S}^{d_\xi-1}$, $\xi_I = \bar{\xi} + AI$, and $\Pi := \{\lambda \geq 0 : G^\top \lambda = b\}$. Consider a proposal distribution $q(dJ | I) = \nu(dJ)$, where ν is the uniform probability measure on $\mathcal{I}$, and the score $\widehat{Q}_x(I) = Q(x, \xi_I)$ for all $I \in \mathcal{I}$. Then, $A_x^{\text{MH}}(I, J) \geq e^{-\Delta/w}$ and $A_x^{\text{B}}(I, J) \geq (1 + e^{\Delta/w})^{-1}$ for all $I, J \in \mathcal{I}$ (recall that $\Delta = \sup_{x \in \mathcal{X}} \sup_{i, j \in \mathcal{I}} |Q(x, \xi^{(i)}) - Q(x, \xi^{(j)})|$ and $\Delta < \infty$ by Assumption 2). It follows that

$$P_x(I, dJ) \geq q(dJ | I) A_x^{\text{MH}}(I, dJ) \geq e^{-\Delta/w} \nu(dJ) \quad (\text{MH})$$

$$P_x(I, dJ) \geq q(dJ | I) A_x^{\text{B}}(I, dJ) \geq (1 + e^{\Delta/w})^{-1} \nu(dJ). \quad (\text{Barker})$$

Then, Doeblin minorization (see, e.g., Lemma 4.3.13 of Cappé et al. 2005) implies that $\eta_{\text{Dob}}(P_x) \leq 1 - e^{-\Delta/w}$ under the MH acceptance and $\eta_{\text{Dob}}(P_x) \leq 1 - (1 + e^{\Delta/w})^{-1}$ under the Barker acceptance.

Alternatively, we can bound $\eta_{\text{Dob}}(P_x)$ through the Lipschitz continuity of $\widehat{Q}_x$. Specifically, recall that $Q(x, \cdot)$ is L'_ξ -Lipschitz continuous in ξ (see Lemma 1). Then, $\widehat{Q}_x(\cdot)$ is $(L'_\xi \|A\|_{2 \rightarrow 2})$ -Lipschitz continuous in I . It follows that, for all $I, J \in \mathcal{I}$,

$$\begin{aligned} A_x^{\text{MH}}(I, J) &= \min \left\{ 1, \frac{(dp_x^w/d\nu)(J)}{(dp_x^w/d\nu)(I)} \right\} \geq \frac{(dp_x^w/d\nu)(J)}{\|dp_x^w/d\nu\|_\infty} \\ A_x^{\text{B}}(I, J) &= \frac{(dp_x^w/d\nu)(J)}{(dp_x^w/d\nu)(I) + (dp_x^w/d\nu)(J)} \geq \frac{(dp_x^w/d\nu)(J)}{2\|dp_x^w/d\nu\|_\infty}. \end{aligned}$$

Hence, for every measurable set $B \subseteq \mathcal{I}$,

$$P_x(I, B) \geq \int_B A_x^{\text{MH}}(I, J) \nu(dJ) \geq \left\| \frac{dp_x^w}{d\nu} \right\|_\infty^{-1} \int_B \frac{dp_x^w}{d\nu}(J) \nu(dJ) = \left\| \frac{dp_x^w}{d\nu} \right\|_\infty^{-1} p_x^w(B) \quad (\text{MH})$$

$$P_x(I, B) \geq \int_B A_x^{\text{B}}(I, J) \nu(dJ) \geq \frac{1}{2} \left\| \frac{dp_x^w}{d\nu} \right\|_\infty^{-1} \int_B \frac{dp_x^w}{d\nu}(J) \nu(dJ) = \frac{1}{2} \left\| \frac{dp_x^w}{d\nu} \right\|_\infty^{-1} p_x^w(B). \quad (\text{Barker})$$

Finally, it can be shown that $\|dp_x^w/d\nu\|_\infty^{-1} \geq C_{d_\xi}(w/(L_\xi \|A\|_{2 \rightarrow 2}))^{d_\xi-1}$ for a constant $C_{d_\xi} > 0$, whose value depends on d_ξ only (see a proof and the specific value of C_{d_ξ} in Appendix EC.4.8). Therefore, Doeblin minorization (see, e.g., Lemma 4.3.13 of Cappé et al. 2005) implies that $\eta_{\text{Dob}}(P_x) \leq 1 - \max \{e^{-\Delta/w}, C_{d_\xi}(w/(L_\xi \|A\|_{2 \rightarrow 2}))^{d_\xi-1}\}$ under the MH acceptance and $\eta_{\text{Dob}}(P_x) \leq 1 - \max \{e^{-\Delta/w}, C_{d_\xi}/2 \cdot (w/(L_\xi \|A\|_{2 \rightarrow 2}))^{d_\xi-1}\}$ under the Barker acceptance. $\square$

Our target is to show that the soft-separation oracle discovers adversarial scenarios as the incumbent sequence and the ensuing Markov kernels co-evolve. To this end, we next make the notion of adversary concrete.

DEFINITION 2 ($\widetilde{\Delta}$ -ADVERSARIAL SCENARIO SET). We define $Q^*(x) := \sup_{I \in \mathcal{I}} Q(x, \xi_I)$ and, for $\widetilde{\Delta} \geq 0$,

$$I_{\widetilde{\Delta}}(x) := \{I \in \mathcal{I} : Q(x, \xi_I) \geq Q^*(x) - \widetilde{\Delta}\}.$$

In addition, with $\widehat{Q}^*(x) := \sup_{I \in \mathcal{I}} \widehat{Q}_x(I)$, we define

$$\widehat{I}_{\widetilde{\Delta}}(x) := \{I \in \mathcal{I} : \widehat{Q}_x(I) \geq \widehat{Q}^*(x) - \widetilde{\Delta}\}.$$

Intuitively, $I_{\widetilde{\Delta}}(x)$ consists of scenarios that are close to the worst-case one with respect to the (true) recourse function and $\widehat{I}_{\widetilde{\Delta}}(x)$ refers to its counterpart with respect to the (approximate) score function. The main result of this section shows that, after sufficiently long co-evolution, the soft-separation oracle identifies $\widetilde{\Delta}$ -adversarial scenarios with high probability.

THEOREM 2 (High-probability identification for general $\mathcal{I}$). *Under Assumptions 5–7, consider the incumbent sequence $\{x_n : n \in \mathbb{N}_+\}$ on a measurable space $(\mathcal{X}, \mathcal{B}_\mathcal{X})$ with respect to the natural filtration $\{\mathcal{F}_n : n \in \mathbb{N}_+\}$ generated by $(\mathcal{I}_m, X_m)_{m \leq n}$:*

$$X_0, I_0 \rightarrow I_1 \rightarrow X_1 \rightarrow I_2 \rightarrow X_2 \rightarrow \dots$$

Conditional on $\mathcal{F}_n$, the next state $\mathcal{I}_{n+1}$ is sampled from $\mathcal{I}_n$ through the one-step Markov kernel P_{x_n} , which admits an invariant distribution $p_{x_n}^w$ on $\mathcal{I}$.

Fix $p > 0$, $\tilde{\Delta} > 2\varepsilon$, and suppose that the incumbent movement satisfies $\|X_n - X_{n-1}\| \leq c\alpha_n$ almost surely, where $\{\alpha_n : n \in \mathbb{N}_+\}$ satisfies $\alpha_n \leq C_\alpha n^{-p}$ for some $C_\alpha > 0$. Then, for every sufficiently large n such that $f(n) := \lceil (p+1) \log n / |\log \bar{\eta}| \rceil < n/2$ and $2^p c C_\alpha f(n) n^{-p} \leq (\tilde{\Delta} - 2\varepsilon)/(4L'_Q)$, it holds that

$$\mathbb{P}\left(I_n \notin I_{\tilde{\Delta}}(X_{n-1})\right) \leq \bar{\eta}^{f(n)} + \frac{2^p c L_P C_\alpha}{1 - \bar{\eta}} f(n) n^{-p} + \sup_{x \in \mathcal{X}} \left\{ \frac{\nu\left(\widehat{I}_{(\tilde{\Delta}-2\varepsilon)/2}(x)^c\right)}{\nu\left(\widehat{I}_{(\tilde{\Delta}-2\varepsilon)/4}(x)\right)} \right\} \cdot \exp\left\{-\frac{\tilde{\Delta} - 2\varepsilon}{4w}\right\}.$$

In particular, the first two terms on the right-hand side scale in order $O((L_P/(1 - \bar{\eta}))n^{-p} \log n)$.

In Theorem 2, the first two terms of the probability bound for not identifying a $\tilde{\Delta}$ -adversarial scenario decay as n increases. For example, if we choose $p = 1$, then the step size $\alpha_n \propto n^{-1}$ implies a decay in the order $O(n^{-1} \log n)$. In addition, the final term evaluates the impact of the approximate score function $\widehat{Q}_x$ on the identification probability. We notice that this term quickly decreases as $(\tilde{\Delta} - 2\varepsilon)$ increases. We elaborate this in the following example, whose proof can be found in Appendix EC.4.10.

EXAMPLE 10. Suppose that $\mathcal{I} = \Xi \subseteq \mathbb{R}^{d_\xi}$ with $\xi_I = I$, Ξ is convex, full-dimensional, and compact with a diameter $D_{\mathcal{I}} := \text{diam}(\Xi) < \infty$, and $\widehat{Q}_x(I) = Q(x, \xi_I)$ with $\varepsilon = 0$. Take ν as the normalized Lebesgue measure on Ξ . Then, for any $\tilde{\Delta} > 0$, it holds that

$$\sup_{x \in \mathcal{X}} \left\{ \frac{\nu\left(\widehat{I}_{\tilde{\Delta}/2}(x)^c\right)}{\nu\left(\widehat{I}_{\tilde{\Delta}/4}(x)\right)} \right\} \leq \max\left\{1, \frac{4L'_\xi D_{\mathcal{I}}}{\tilde{\Delta}}\right\}^{d_\xi}$$

and therefore the probability bound of Theorem 2 is in the order $O(n^{-1} \log n + \exp\{-\tilde{\Delta}/(4w)\})$. In other words, the longer we run the Markov chain and the larger $\tilde{\Delta}$ is, the soft-separation oracle becomes more powerful for identifying adversarial scenarios. $\square$

We turn to the probability bound for identifying adversarial scenarios in a finite state space $\mathcal{I}$. The following theorem takes advantage of the finiteness of $\mathcal{I}$ and a uniform gap between maximal and other recourse costs to yield a geometric decay.

THEOREM 3 (High-probability identification for finite $\mathcal{I}$). *Under the assumptions of Theorem 2, suppose that $\mathcal{I}$ is finite. Let ν be the uniform distribution on $\mathcal{I}$, $I^*(x) := \{i \in \mathcal{I} : Q(x, \xi^{(i)}) = \max_{I \in \mathcal{I}} Q(x, \xi^{(I)})\}$, $\tilde{\Delta}_n := \max_{I \in \mathcal{I}} Q(x_n, \xi^{(I)}) - \max_{j \notin I^*(x_n)} Q(x_n, \xi^{(j)})$, and suppose that there exists a $\Delta_\star > 4\varepsilon$ such that $\Delta_\star \leq \tilde{\Delta}_n$ for all $n \in \mathbb{N}_+$. Then, for all $n \geq 2$, it holds that*

$$\mathbb{P}\left(I_n \notin I^*(X_{n-1})\right) \leq \bar{\eta}^{n-1} \mathbb{P}(I_1 \notin I^*(X_0)) + \frac{2c\bar{\eta}L'_Q}{\Delta_\star - 4\varepsilon} \sum_{t=1}^{n-1} \bar{\eta}^{n-1-t} \alpha_t + \frac{1 - \bar{\eta}^{n-1}}{1 - \bar{\eta}} (|\mathcal{I}| - 1) e^{-(\Delta_\star - 2\varepsilon)/w}.$$

In addition, for all $n \geq n_0 := \max \left\{ 2, 1 + \left\lceil \left(\frac{2cC_\alpha L'_Q}{\Delta_\star - 4\varepsilon} \right)^{1/p} \right\rceil \right\}$, it holds that

$$\mathbb{P}(I_n \notin I^\star(X_{n-1})) \leq \bar{\eta}^{n-n_0} + \frac{1 - \bar{\eta}^{n-1}}{1 - \bar{\eta}} (|\mathcal{I}| - 1) e^{-(\Delta_\star - 2\varepsilon)/w}.$$

In Algorithm 1, one can apply the soft separation to establish $\underline{F}_n := \min_{x \in \mathcal{X}} \{f(x) + \max_{m \in [n]} Q(x, \xi_{I_m})\}$ and $\bar{F}_n := f(x_{n-1}) + Q(x_{n-1}, \xi_{I_n}) + \tilde{\Delta}$ in each iteration n . We notice that $\underline{F}_n$ is a deterministically valid lower bound of $F^\star$ while, by Theorems 2–3, $\bar{F}_n$ is a high-confidence upper bound for $F^\star$. Hence, one can design an early termination criterion for Algorithm 1 based on $\underline{F}_n$ and $\bar{F}_n$. Finally, we extend Theorem 1 to a general state space $\mathcal{I}$ through the soft-separation oracle.

THEOREM 4. Fix $w > 0$, let P_x be the Markov kernel induced by the uniform proposal distribution on $\mathcal{I}$ and the MH acceptance, and set the step sizes as $\alpha_n = \alpha_0/\sqrt{n}$ in Algorithm 1 for an $\alpha_0 > 0$. Then, under Assumptions 1, 6, and 7, it holds that

$$\mathbb{E}[F_w^\nu(\bar{x}_n)] - \min_{x \in \mathcal{X}} F_w^\nu(x) = \begin{cases} O \left[(C_\star + C_\star^2/w) \frac{\log n}{\sqrt{n}} \right] & \text{for Step-size Average} \\ O \left[(C_\star + C_\star^2/w) \frac{1}{\sqrt{n}} \right] & \text{for Linear Average} \end{cases}$$

where $C_\star = (1 + \bar{\eta})/(1 - \bar{\eta})$.

4. Mixed-Integer Here-and-Now Decisions

We extend the application of the soft separation to (ARO) with mixed-integer decision variables, i.e., $\mathcal{X} \subseteq \mathbb{Z}^p \times \mathbb{R}^{d_x - p}$. For ease of exposition, we write $x := (x_z, x_c)$, where x_z and x_c denote the integer and continuous entries of x , respectively. We shall develop a branch-and-cut method that generates valid cuts for the robust objective through the soft-separation oracle from Section 3. To avoid clutter, we restrict ourselves to the *linear* recourse, i.e., we make Assumptions 2–3 and work with a finite state space $\mathcal{I}$.

We rewrite (ARO) as an epigraphic formulation

$$\min_{x \in \mathcal{X}, \theta \in \mathbb{R}} f(x) + \theta \quad \text{s.t.} \quad \theta \geq Q(x, \xi^{(i)}) \quad \forall i \in \mathcal{I}.$$

For any scenario $i \in \mathcal{I}$, strong duality gives $Q(x, \xi^{(i)}) = \max_{\pi \in \Pi} \pi^\top (h - Ex - U\xi^{(i)})$. Hence, for all $\bar{\pi}_i \in \arg \max_{\pi \in \Pi} \pi^\top (h - E\bar{x} - U\xi^{(i)})$ with respect to any $\bar{x} \in \mathcal{X}$, the (dual Benders') cut

$$\theta \geq \bar{\pi}_i^\top (h - Ex - U\xi^{(i)}) \tag{4}$$

is valid for the robust epigraph. We apply the soft separation to identify adversarial scenarios $\xi^{(i)}$ and then apply the corresponding Benders' cut to solve (ARO) through branch-and-cut.

4.1. Applying the soft-separation oracle in branch-and-cut

At an integer-feasible incumbent $(\bar{x}, \bar{\theta})$, we apply the soft separation established in Section 3 to return a (high-confidence) adversarial scenario and generate a deterministically valid cut. The main technical issue, however, is that the probabilistic guarantee of Theorem 3 requires the incumbent sequence to evolve gradually. In particular, the theorem assumes that consecutive incumbents satisfy $\|x_n - x_{n-1}\| \leq \alpha_n$ for a decaying sequence $\alpha_n \sim n^{-p}$. A generic branch-and-bound incumbent sequence need not satisfy this requirement, because consecutive incumbents may have different integer components and are far apart. To make the tracking analysis applicable, we maintain a Markov chain for each newly discovered integer assignment x_z and will continue using the same chain for all future incumbents x that share the same integer assignment.

Specifically, for each feasible solution $\bar{x} \equiv (\bar{x}_z, \bar{x}_c) \in \mathbb{Z}^p \times \mathbb{R}^{d_x - p}$ encountered in the branch-and-bound tree, we maintain **state** $(\bar{x}_z)$, **ref** $(\bar{x}_z)$, and $\kappa(\bar{x}_z)$. Here, **state** $(\bar{x}_z) \in \mathcal{I}$ records the most recently sampled adversarial scenario associated with $\bar{x}_z$, **ref** $(\bar{x}_z) \in \mathbb{R}^{d_c}$ records the continuous entries of the corresponding incumbent, and $\kappa(\bar{x}_z) \in \mathbb{N}_+$ records the number of previous oracle calls associated with $\bar{x}_z$. When $\bar{x}$ is encountered, we compare $\bar{x}_c$ with the previously recorded **ref** $(\bar{x}_z)$. If

$$\|\bar{x}_c - \mathbf{ref}(\bar{x}_z)\| \leq \frac{c_0}{\max\{\kappa(\bar{x}_z), 1\}^p}, \quad (5)$$

then $\bar{x}$ stays within the neighborhood designated in Theorem 3 with a width $\alpha_n = c_0 \max\{\kappa(\bar{x}_z), 1\}^{-p}$. The oracle therefore performs only one transition under $P_{\bar{x}}(\mathbf{state}(\bar{x}_z), \cdot)$. On the other hand, if (5) fails to hold, then the target distribution $p_{\bar{x}}$ may have changed too much for the one-step tracking argument to apply. In this case, the oracle performs $g(\kappa(\bar{x}_z)) \geq 1$ transitions under the fixed kernel $P_{\bar{x}}$ before returning an adversarial scenario. This yields an adaptive Markov chain, where the incumbent evolves at a different pace from the state, as opposed to a co-evolution as in Theorem 3. We therefore extend Theorem 3 to find $g(\kappa(\bar{x}_z))$.

COROLLARY 1 (Identification under adaptive transition steps). *Suppose the assumptions of Theorem 3 hold, with $\Delta_\star > 4\varepsilon$. Consider the Markov chain associated with any fixed integer decision $\bar{x}_z$, let $\kappa = \kappa(\bar{x}_z) \geq 1$, and we perform one transition whenever (5) holds and $g(\kappa)$ transitions otherwise, where*

$$g(\kappa) := \max \left\{ 1, \left\lceil \frac{p \log(\max\{\kappa, 1\}) - \log\left(\frac{2c_0\bar{\eta}L'_Q}{\Delta_\star - 4\varepsilon}\right)}{|\log \bar{\eta}|} \right\rceil \right\} = O(\log(\kappa)). \quad (6)$$

Then, it holds that

$$\mathbb{P}\left(\mathbf{state}(\bar{x}_z) \notin I^\star(\mathbf{ref}(\bar{x}_z))\right) \leq \bar{\eta}^{\kappa-1} + \frac{2c_0\bar{\eta}L'_Q}{\Delta_\star - 4\varepsilon} \sum_{t=1}^{\kappa-1} \bar{\eta}^{\kappa-1-t} t^{-p} + \frac{|\mathcal{I}| - 1}{1 - \bar{\eta}} (1 - \bar{\eta}^{\kappa-1}) \exp\left\{-\frac{\Delta_\star - 2\varepsilon}{w}\right\}.$$

In particular, with constant $C_{\bar{\eta},p} := \frac{1}{1-\bar{\eta}} [2^p + \bar{\eta}^{-1/2} (2p/(e|\log \bar{\eta}|))^p]$, the right-hand side of the last inequality is bounded from above by

$$\widehat{\delta}(\kappa) := \frac{|\mathcal{I}| - 1}{1 - \bar{\eta}} \exp\left\{-\frac{\Delta_\star - 2\varepsilon}{w}\right\} + \bar{\eta}^{\kappa-1} + \frac{2c_0\bar{\eta}L_{Q'}}{\Delta_\star - 4\varepsilon} C_{\bar{\eta},p} \kappa^{-p}.$$

We summarize this adaptive variant of the soft-separation oracle in Algorithm 2.

Algorithm 2 Adaptive Soft-Separation Oracle

Input: Incumbent $(\bar{x}, \bar{\theta})$ with $\bar{x} = (\bar{x}_z, \bar{x}_c)$; memory maps **state**($\cdot$), **ref**($\cdot$), $\kappa(\cdot)$

Output: State I , valid cut, and violation indicator

- 1: **if** $\bar{x}_z$ has not been observed before **then**
 - 2: Initialize **state**($\bar{x}_z$) $\in \mathcal{I}$, **ref**($\bar{x}_z$) $\leftarrow \bar{x}_c$, and $\kappa(\bar{x}_z) \leftarrow 0$
 - 3: $m \leftarrow \begin{cases} 1, & \text{if } \kappa(\bar{x}_z) = 0 \text{ or } \|\bar{x}_c - \mathbf{ref}(\bar{x}_z)\| \leq c_0 \max\{\kappa(\bar{x}_z), 1\}^{-p} \\ g(\kappa(\bar{x}_z)), & \text{otherwise, through (6)} \end{cases}$
 - 4: Sample $I \sim P_{\bar{x}}^m(\mathbf{state}(\bar{x}_z), \cdot)$
 - 5: Define a valid cut through (4) with $i = I$
 - 6: $\text{violated} \leftarrow \mathbf{1}\{Q(\bar{x}, \xi^{(I)}) > \bar{\theta}\}$
 - 7: **state**($\bar{x}_z$) $\leftarrow I$, **ref**($\bar{x}_z$) $\leftarrow \bar{x}_c$, and $\kappa(\bar{x}_z) \leftarrow \kappa(\bar{x}_z) + 1$
 - 8: **return** I , the valid cut, and violated
-

4.2. Branch-and-cut with certification

We incorporate the adaptive soft-separation oracle (Algorithm 2) in a branch-and-cut algorithm for solving (ARO) with mixed-integer x . When implementing this in Gurobi or CPLEX, we can embed Algorithm 2 via the `lazy callback`. Let $\mathcal{C} \subseteq \mathcal{I} \times \Pi$ denote the pool of cuts produced by Algorithm 2. This produces a relaxation of (ARO),

$$L(\mathcal{C}) := \min_{x \in \bar{\mathcal{X}}, \theta \in \mathbb{R}} \{f(x) + \theta : \theta \geq \pi^\top (h - Ex - U\xi^{(i)}), (i, \pi) \in \mathcal{C}\}, \quad (7)$$

where $\bar{\mathcal{X}}$ denotes a relaxation of $\mathcal{X}$ produced by the branch-and-cut method. Because every cut in $\mathcal{C}$ is deterministically valid, $L(\mathcal{C})$ is a lower bound for $F^\star$, i.e., $L(\mathcal{C}) \leq F^\star$. On the other hand, branch-and-cut terminates when the upper bound

$$\widehat{U} := f(\bar{x}) + \bar{\theta}$$

associated with an integer incumbent $(\bar{x}, \bar{\theta})$ is close enough to $L(\mathcal{C})$, e.g., when $\widehat{U} - L(\mathcal{C}) \leq \varepsilon_{\text{opt}}$ with a prespecified optimality gap ε_{opt} . However, $\widehat{U}$ is not necessarily a valid upper bound for $F^\star$ because $\bar{\theta}$ may underestimate $Q^\star(\bar{x}) \equiv \max_{\xi \in \Xi} Q(\bar{x}, \xi)$. Consequently, the branch-and-cut method

alone may not certify an ε_{opt} -optimal solution to (ARO). To this end, we propose to embed a certification procedure within the branch-and-cut to guarantee that $(\bar{x}, \bar{\theta})$ is an ε_{opt} -optimal solution with probability at least $1 - \varepsilon_{\text{false}}$, where $\varepsilon_{\text{false}}$ denotes a threshold for falsely identifying an optimal solution.

This procedure solves (ARO) in a single branch-and-bound tree and, for each integer incumbent $\bar{x}$ encountered, employs Algorithm 2 to (i) find a Benders' cut (4) and (ii) certify a high-confidence upper bound $\hat{U}$. Whenever a violated cut is found, we add it to $\mathcal{C}$ and continue the branch-and-bound. Otherwise, we repeat calling Algorithm 2 at $\bar{x}$ for a sufficiently many times to certify $\hat{U}$.

First, to find a Benders' cut, we make one call to Algorithm 2 when we encounter an integer incumbent $(\bar{x}, \bar{\theta})$ in the branch-and-bound tree. If no violated cut is found, we record $k = \kappa(\bar{x}_z)$ after this call. Consequently, k includes all previous calls associated with $\bar{x}_z$, including earlier certification calls specified next.

Second, to certify upper bounds, we assign a label $j = 1, 2, \dots$ to each Markov chain before its first sampling and define

$$\rho := \bar{\eta} + (|\mathcal{I}| - 1) \exp\left\{-\frac{\Delta_\star - 2\epsilon}{w}\right\}, \quad r := \frac{j(j+1)k(k+1)}{\epsilon_{\text{false}}} \min\{1, \hat{\delta}(k)\}. \quad (8)$$

At each integer incumbent $(\bar{x}, \bar{\theta})$, we perform $\max\{0, \lceil \log_{1/\rho}(r) \rceil\}$ additional calls to Algorithm 2. Note that, because $\bar{x}_c = \mathbf{ref}(\bar{x}_z)$ is fixed throughout these calls, each additional call uses one transition of $P_{\bar{x}}$ (i.e., $m = 1$ in Algorithm 2). If no cuts are found violated in these calls, then we certify $\hat{U} = f(\bar{x}) + \bar{\theta}$ as an upper bound.

We now analyze the probability of incorrectly certifying an upper bound $\hat{U}$ with $\bar{\theta} < Q^\star(\bar{x})$. Recall, from the proof of Theorem 3, that

$$P_x(I, I^\star(x)) \geq 1 - \rho, \quad \forall x \in \mathcal{X}, I \notin I^\star(x), \quad (9)$$

where we choose a temperature parameter $w \leq (\Delta_\star - 2\epsilon) / \log(2(|\mathcal{I}| - 1)/(1 - \bar{\eta}))$ to ensure that $\rho \leq (1 + \bar{\eta})/2 < 1$. Because the first call of Algorithm 2 (in the ‘‘cut’’ phase) does not find a Benders' cut, the state sampled therein lies outside of $I^\star(\bar{x})$. Consequently, inequality (9) designates that each additional call of Algorithm 2 (in the ‘‘certification’’ phase) fails to find $\bar{\theta} < Q^\star(\bar{x})$ with probability at most ρ . Combining this conditional bound with the marginal bound of Corollary 1 gives, for each fixed (j, k) ,

$$\mathbb{P}\left\{\begin{array}{l} \text{certification at } (j, k) \\ \text{with } Q^\star(\bar{x}) > \bar{\theta} \end{array}\right\} \leq \min\{1, \hat{\delta}(k)\} \rho^{\max\{0, \lceil \log_{1/\rho} r \rceil\}} \leq \frac{\epsilon_{\text{false}}}{j(j+1)k(k+1)}.$$

Note that each call of Algorithm 2 increases $\kappa(\bar{x}_z)$ by one, so every certification is associated with a unique (j, k) pair. Then, a union bound gives

$$\mathbb{P}\left\{\begin{array}{l} \text{certifying any } (\bar{x}, \bar{\theta}) \\ \text{with } Q^\star(\bar{x}) > \bar{\theta} \end{array}\right\} \leq \epsilon_{\text{false}} \sum_{j=1}^{\infty} \sum_{k=1}^{\infty} \frac{1}{j(j+1)k(k+1)} = \epsilon_{\text{false}}.$$

Finally, we notice that each certification concludes with either an acceptance of $\widehat{U}$ or a violated cut, and does so within finite iterations because $\log_{1/\rho}(r) < \infty$. This, together with the finiteness of cuts and the finite size of the branch-and-bound tree, guarantees that the proposed branch-and-cut method is terminated finitely. We summarize the proposed certification procedure in Algorithm 3 and its probabilistic guarantee in the following theorem.

THEOREM 5. *Algorithm 3 terminates finitely almost surely. In addition, with probability at least $1 - \varepsilon_{\text{false}}$, its returned solution (x^*, θ^*) satisfies $0 \leq (f(x^*) + \theta^*) - F^* \leq \varepsilon_{\text{opt}}$.*

Algorithm 3 Branch-and-Cut with Certification

Input: Cut pool $\mathcal{C}$; maps **state**, **ref**, κ ; budget $\varepsilon_{\text{false}}$; temperature parameter w ;

Output: Solution x^*

- 1: Solve the root-node relaxation using Algorithm 1 and warm-start $\mathcal{C} \leftarrow \{(I_{n+1}, \pi^*(x_n, \xi^{(I_{n+1})}))\}_n$
 - 2: Set $j \leftarrow 0$, $\widehat{U} \leftarrow +\infty$, $L(\mathcal{C}) \leftarrow$ root-node optimal value
 - 3: **while** $\widehat{U} - L(\mathcal{C}) > \varepsilon_{\text{opt}}$ **do** solve (ARO) using branch-and-bound
 - 4: **for** each integer incumbent $(\bar{x}, \bar{\theta})$ encountered **do** ▷ lazy callback
 - 5: **if** the chain for x_z has no label **then** assign label j to this chain and $j \leftarrow j + 1$
 - 6: Call Algorithm 2 at $(\bar{x}, \bar{\theta})$, set $k \leftarrow \kappa(\bar{x}_z)$, and set ρ, r through (8) ▷ “cut” phase
 - 7: **while** true **do**
 - 8: **if** a violated cut is returned **then** add it to $\mathcal{C}$ and **break**
 - 9: **else if** $r \leq 1$ **then** ▷ “certification” phase
 - 10: **if** $f(\bar{x}) + \bar{\theta} < \widehat{U}$ **then** update $\widehat{U} \leftarrow f(\bar{x}) + \bar{\theta}$ and $x^* \leftarrow \bar{x}$
 - 11: Certify $\widehat{U}$ and **break**
 - 12: **else** call Algorithm 2 at $(\bar{x}, \bar{\theta})$ and update $r \leftarrow \rho r$
 - 13: Update $L(\mathcal{C}) \leftarrow$ the global lower bound of the branch-and-bound tree
 - 14: **return** x^*
-

5. Numerical Experiments

We evaluate the effectiveness and scalability of the proposed Algorithms 1 and 3 numerically. In particular, Section 5.1 applies Algorithm 1 to solving (ARO) models with continuous x , and Section 5.2 applies Algorithm 3 to address mixed-integer x . In both cases, we compare the proposed algorithm to C&CG and ADR under various instance settings. Unless otherwise stated, each configuration is repeated over 20 independent replications, and the reported curves show the median trajectory together with the 25–75% interquartile band. All experiments were conducted in Python

on a MacBook Pro equipped with an Apple M3 Pro Chip and 18 GB RAM. All instances and code are available online at <https://github.com/xuqy2002/SoftSeparationAR0>.

As benchmarks, we implement C&CG of Zeng and Zhao (2013) to solve the instances to global optimum and the ADR approximation of Ben-Tal et al. (2004). To compare fairly with C&CG, in the case of an ellipsoidal uncertainty set, we follow Electronic Companion A-4.1 in Zeng and Zhao (2013) to solve its subproblems based on the KKT condition of the transportation linear program and the big- M constants suggested therein; in the case of scenario uncertainty set, we compute $\max_{\xi \in \Xi} Q(x, \xi)$ through enumeration in C&CG, which we find significantly faster than computing it through the KKT condition plus an SOS1 constraint among the ξ -scenarios; and in the case of a budget uncertainty set, we follow Electronic Companion A-4.2 in Zeng and Zhao (2013) to solve its subproblems based on strong duality and the big- M constants in Gabrel et al. (2014), as Zeng and Zhao (2013) suggest. Other implementation details can be found in Appendix EC.3.

5.1. Continuous Here-and-Now Decisions

We consider the adaptive robust location-transportation problem as in Zeng and Zhao (2013). The here-and-now decision allocates continuous capacity across m supply locations at a linear investment cost. After the uncertain demand vector $\xi \in \mathbb{R}^n$ is realized, a wait-and-see transportation linear program ships goods to n demand locations subject to the installed capacities, with any unfulfilled demand served by a high-cost emergency source. The demand at location $j \in [n]$ is modeled by $d_j(\xi) = \bar{d}_j + \hat{d}_j \xi_j$, where the primitive uncertainty ξ belongs with an uncertainty set.

We generate random instances as follows. We set $m = n$ (i.e., equal number of supply and demand nodes), and sample the integer-valued capacity-investment costs, transportation costs, emergency-supply costs, and nominal demands $\bar{d}_j$ uniformly and independently from the intervals $[10, 30]$, $[20, 40]$, $[100, 200]$, and $[200, 300]$, respectively. The demand sensitivity $\hat{d}_j$ is a uniform random fraction $[0.1, 0.5]$ of $\bar{d}_j$. We consider three types of uncertainty sets: firstly, an ellipsoidal uncertainty set

$$\Xi := \left\{ \bar{\xi} + R \text{diag}(a_1, \dots, a_n) u : \|u\|_2 \leq 1 \right\}, \quad R^T R = \text{Id}, \quad (\text{Ellipsoidal})$$

where $\bar{\xi} = 0.5 \cdot \mathbf{1}$, the semi-axis lengths a_j are sampled independently from $\text{Unif}[0.10, 0.50]$, and R is an independent random rotation matrix drawn from the Haar distribution on the rotation group $\text{SO}(n)$; and secondly, a scenario uncertainty set consisting of finite scenarios of ξ , which are drawn uniformly from the boundary of the ellipsoidal uncertainty set. Consequently, these two uncertainty sets produce similar F^* . Thirdly, a budget uncertainty set is defined through

$$\Xi := \left\{ \xi \in [0, 1] : \mathbf{1}^T \xi \leq n\Gamma \right\}, \quad (\text{Budget})$$

where $n\Gamma$ is the so-called “budget of uncertainty.”

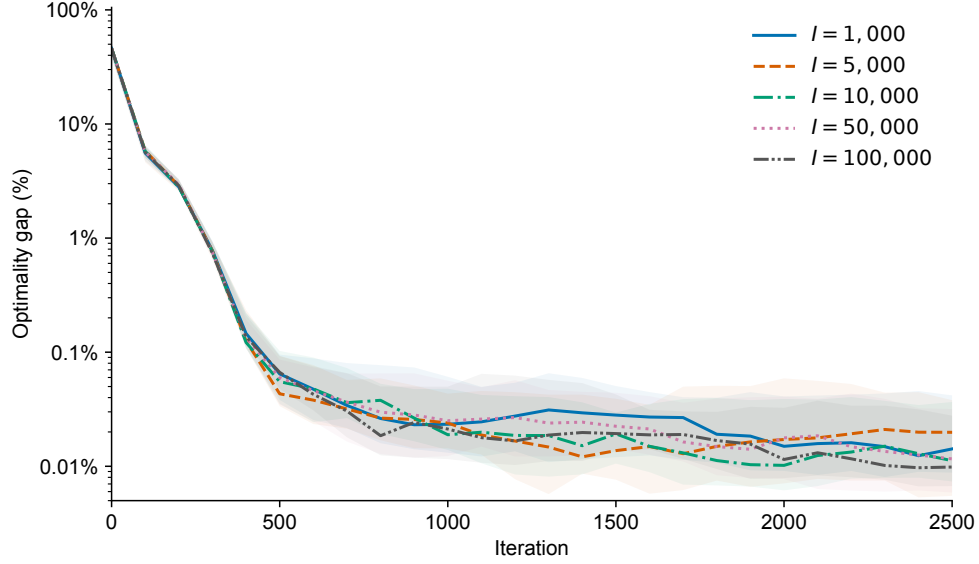


Figure 1 Convergence of Algorithm 1 under various scenario sizes $|\mathcal{I}|$ with uniform proposal and $w = 0.001$. Each selected solution is reevaluated using the exact separation. The solid curves show medians over 20 replications, and the shaded regions show the 25–75% interquartile bands.

5.1.1. Hyperparameter Configuration We first examine the convergence of Algorithm 1 under various choices of $|\mathcal{I}|$ (the size of the scenario uncertainty set), w (temperature parameter for the Mellowmax approximation), and ν (the proposal distribution on $\mathcal{I}$). To this end, we evaluate the robust objective function value of each decision incumbent using exact separation, and report the evolution of the ensuing relative optimality gap when compared to the true optimal value. In all evaluations, we use the linear average as the decision incumbent.

Convergence under varying scenario sizes. We first examine the convergence of Algorithm 1 with $|\mathcal{I}| \in \{1000, 5000, 10000, 50000, 100000\}$, while fixing $w = 0.001$ and ν to be uniform on $\mathcal{I}$. Figure 1 reports the convergence trajectories. We observe that Algorithm 1 displays broadly similar convergence behavior across all tested scenario sizes. In each trajectory, the relative optimality gap decreases sharply in the initial iterations and then stabilizes around the interval $[0.01\%, 0.1\%]$. In particular, increasing $|\mathcal{I}|$ from 1,000 to 100,000 does not lead to a systematic deterioration in either convergence speed or the attained gap over the iteration horizon considered. This demonstrates that the convergence of Algorithm 1 has a mild reliance on the scenario size.

Sensitivity to the temperature parameter w . We next study the convergence with w varying among $\{1, 5, 10, 50, 100\}$, while fixing ν to be uniform on $\mathcal{I}$ and $|\mathcal{I}| = 1,000$. From Figure 2, we observe that the optimality gap displays a phase change, with $w \in \{50, 100\}$ resulting in significant gaps between 0.1% and 1%, and lower temperature such as $w \leq 10$ producing much smaller gaps in $[0.01\%, 0.1\%]$. In addition, the gap generally improves as the w -value decreases. This reflects the

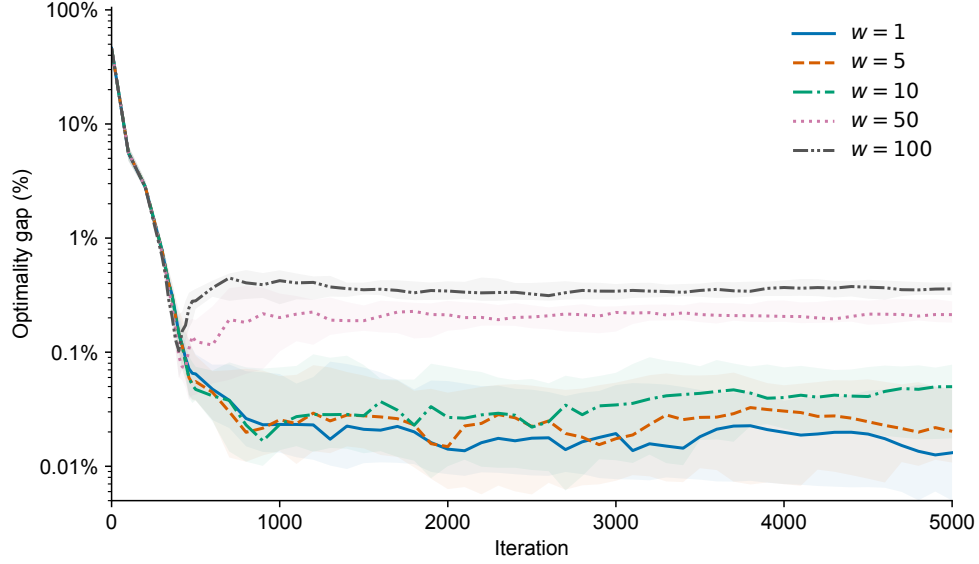


Figure 2 Convergence of Algorithm 1 under various w with uniform proposal and $|\mathcal{I}| = 1,000$. Each selected solution is reevaluated using the exact separation. The solid curves show medians over 20 replications, and the shaded regions show the 25–75% interquartile bands.

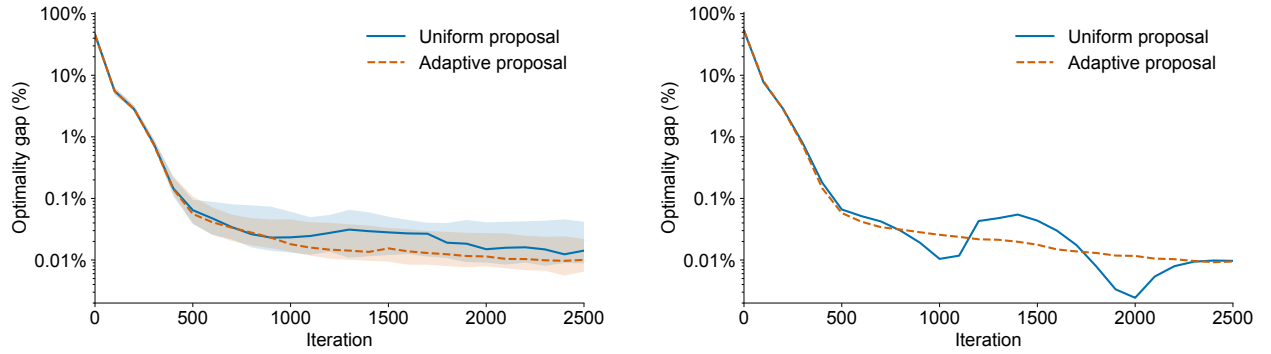


Figure 3 Adaptive proposal versus uniform proposal with $w = 0.001$ and $|\mathcal{I}| = 1,000$. Each selected solution is reevaluated using the exact separation. The left panel reports the median trajectories over 20 independent replications, with the shaded regions indicating the 25–75% interquartile bands. The right panel shows the convergence trajectories from one individual replication.

Mellowmax approximation error in Lemma 2. We note that the error gap therein is uniform in x and so it makes sense that Figure 2 displays a much higher accuracy around optimum. In the remaining experiments, we shall fix $w = 0.001$ for an accurate approximation of (ARO).

Adaptive proposal versus uniform proposal. Finally, we examine the proposal distribution ν . In Figure 3–right, we depict the convergence of Algorithm 1 with the uniform proposal, $w = 0.001$, and $|\mathcal{I}| = 1,000$. The solid curve displays a decaying trend in the optimality gap but also shows rebounds in later iterations (note that the gap is reported in a logarithmic scale). Intuitively, we can attribute these rebounds to the uniform proposal because, in later iterations, a uniform ν ignorant

of $Q(x, \xi)$ may propose less adversarial states which, if accepted, move the decision incumbents away from optimum.

Observing this, we propose an adaptive proposal distribution to account for $Q(x, \xi)$. To this end, we use the subgradient $-U^\top \pi^*(x, i) \in \partial_\xi Q(x, \xi^{(i)})$ at the current state $i \in \mathcal{I}$ to tilt new scenario generation toward a higher $Q(x, \xi)$ value. Specifically, for the ellipsoidal uncertainty set defined in (Ellipsoidal), we draw the primitive random vector u from the von Mises–Fisher (vMF) distribution on $\mathbb{S}^{n-1}$, which has a density $f_{\text{vMF}}(u; \mu, \kappa) \propto \exp(\kappa \mu^\top u)$ with respect to the uniform spherical measure. Here, $\mu := \frac{A^\top (-U^\top \pi^*(x, i))}{\|A^\top (-U^\top \pi^*(x, i))\|}$ with $A := R \text{diag}(a_1, \dots, a_n)$ is the mean direction and $\kappa \geq 0$ controls the concentration around μ . For example, $\kappa = 0$ reduces vMF to the uniform distribution on $\mathbb{S}^{n-1}$ and a larger κ places more probability mass near μ . Then, we propose $\bar{\xi} + Au$. For the scenario uncertainty set, we adapt the proposal distribution to be $q(j|i) \propto \frac{\exp(\kappa \mu^\top u^{(j)})}{|\det A| \|A^{-\top} u^{(j)}\|}$. For the budget uncertainty set, the adaptive proposal admits an efficient sampling procedure through dynamic programming, which we detail in Appendix EC.2.

We compare the adaptive proposal with $\kappa = 10$ and the uniform proposal in Figure 3. We observe that the adaptive proposal produces a more stable convergence (see the right panel) and converges to a smaller optimality gap (see the left panel). This makes sense because vMF tilts the proposal toward adversarial states and reduces ineffective transitions. It can also be explained by the Poisson resolvent constant C_* . Indeed, the adaptive proposal is closer to the ideal sampling distribution p_x than the uniform proposal, leading to a smaller C_* value (see Example 4) and a faster convergence (see Theorems 1 and 4).

5.1.2. Computing Time We report the computing times (all in wall clock seconds) of solving (ARO) instances through Algorithm 1 using $w = 0.001$ and the adaptive proposal distribution.

Table 1 Runtime comparison in instances with *continuous* here-and-now decisions and ellipsoidal uncertainty sets. Entries are mean $\pm$ standard deviation in seconds over five independent replications.

n	Algorithm 1	C&CG	ADR
3	0.13 ± 0.06	0.045 ± 0.016	0.029 ± 0.004
10	0.17 ± 0.07	2.73 ± 0.43	1.33 ± 0.51
30	0.32 ± 0.10	772.78 ± 376.24	138.52 ± 11.39
50	1.07 ± 0.25	> 1800	623.83 ± 107.04
70	2.07 ± 0.55	> 1800	> 1800
100	3.56 ± 0.65	> 1800	> 1800

Ellipsoidal uncertainty set. We first consider the ellipsoidal uncertainty set (Ellipsoidal) with n varying from 3 to 100 and report the computing times in Table 1. For Algorithm 1 and C&CG, we report the time required to attain a relative optimality gap below 0.1% with respect to the true optimal value. For ADR, we report the time required to solve the conservative approximation. We

use a time limit of half an hour. We report “> 1800” if the computing times in all replications exceed the limit, or otherwise truncate the beyond-limit computing times at 1,800s and then report mean $\pm$ standard deviation in seconds.

From Table 1, we observe that Algorithm 1 shows better scalability than C&CG and ADR as the problem dimension n increases. In particular, the computing times of Algorithm 1 grow mildly and remain stable as n increases. On the contrary, those of C&CG and ADR exceed the time limit in all replications once n reaches 70. This demonstrates the scalability of the proposed soft separation approach with problem dimensions.

Table 2 Runtime comparison in instances with *continuous* here-and-now decisions and scenario uncertainty sets. Entries are mean $\pm$ standard deviation in seconds over five independent replications.

$ \mathcal{I} $	$n = 3$			$n = 30$			$n = 100$		
	Algorithm 1	C&CG	ADR	Algorithm 1	C&CG	ADR	Algorithm 1	C&CG	ADR
1,000	0.15 $\pm$ 0.04	0.17 $\pm$ 0.03	0.14 $\pm$ 0.08	0.74 $\pm$ 0.19	1.07 $\pm$ 0.08	> 1800	6.71 $\pm$ 2.25	9.54 $\pm$ 0.76	> 1800
5,000	0.17 $\pm$ 0.06	0.81 $\pm$ 0.32	2.52 $\pm$ 0.14	0.57 $\pm$ 0.23	5.03 $\pm$ 0.43	> 1800	7.87 $\pm$ 0.41	48.94 $\pm$ 3.45	> 1800
10,000	0.21 $\pm$ 0.08	1.46 $\pm$ 0.52	4.85 $\pm$ 0.80	0.67 $\pm$ 0.10	9.92 $\pm$ 0.61	> 1800	8.01 $\pm$ 0.93	94.45 $\pm$ 8.39	> 1800
50,000	0.37 $\pm$ 0.14	5.46 $\pm$ 1.93	13.42 $\pm$ 0.45	1.02 $\pm$ 0.32	49.44 $\pm$ 3.77	> 1800	7.72 $\pm$ 0.27	433.20 $\pm$ 32.83	> 1800
100,000	0.58 $\pm$ 0.21	12.58 $\pm$ 5.35	54.43 $\pm$ 11.83	1.26 $\pm$ 0.25	101.95 $\pm$ 5.93	> 1800	8.29 $\pm$ 0.29	844.09 $\pm$ 94.49	> 1800

Scenario uncertainty set. We next consider the scenario uncertainty set $\Xi := \{\xi^1, \dots, \xi^{|\mathcal{I}|}\}$ with $|\mathcal{I}|$ varying from 1,000 to 100,000 and n varying from 3 to 100. From Table 2, we observe that the runtime of Algorithm 1 increases only moderately in $|\mathcal{I}|$, reflecting the uniform error bound in Lemma 2 and the convergence rate in Theorem 1. In addition, the increase of Algorithm 1 runtime in n is also significantly milder than those of C&CG and ADR. In particular, ADR cannot solve any instance/replication within the time limit once n reaches 30. This observation supports the computational advantage of the soft separator: it avoids the scenario-wise scaling bottleneck, which limits C&CG and ADR in higher-dimensional instances.

Budget uncertainty set. Finally, we consider the budget uncertainty set (**Budget**) with n varying from 30 to 100 and Γ varying between 0.1 and 0.9. From Table 3, we observe that C&CG performs the best when n is small and Γ is near-binary (e.g., when $n = 30, \Gamma = 0.9$). As n increases and Γ moves away from binary values, Algorithm 1 outperforms C&CG and ADR. For example, neither C&CG nor ADR solves any instance or replication within the time limit once n reaches 100, while Algorithm 1 solves all these instances. Similarly, the Algorithm 1 runtime remains more stable than that of C&CG as Γ varies. For example, the C&CG runtime with $\Gamma = 0.5$ (when the budget uncertainty set has more extreme points) is significantly longer than that with $\Gamma = 0.1, 0.9$ (when the number of extreme points is less).

Table 3 Runtime comparison for instances with *continuous* here-and-now decisions and budget uncertainty sets. Times are reported as mean $\pm$ standard deviation in seconds over five independent replications.

n	Γ	Algorithm 1	C&CG	ADR	n	Γ	Algorithm 1	C&CG	ADR
30	0.1	1.46 ± 0.25	1.13 ± 0.61	3.94 ± 0.40	70	0.1	17.79 ± 2.35	> 1800	473.61 ± 65.69
	0.3	1.56 ± 0.37	18.54 ± 16.68	7.78 ± 1.94		0.3	10.88 ± 0.55	> 1800	768.04 ± 188.41
	0.5	2.33 ± 1.15	10.54 ± 10.25	13.67 ± 5.98		0.5	8.02 ± 2.35	> 1800	881.43 ± 209.58
	0.7	2.22 ± 0.36	2.74 ± 2.92	14.31 ± 6.08		0.7	10.80 ± 1.37	> 1800	1027.86 ± 318.44
	0.9	1.90 ± 0.55	0.20 ± 0.04	16.63 ± 4.97		0.9	23.49 ± 5.42	63.93 ± 54.06	1319.22 ± 573.24
50	0.1	8.51 ± 1.54	37.33 ± 16.57	74.58 ± 15.31	100	0.1	38.79 ± 7.96	> 1800	> 1800
	0.3	6.04 ± 0.76	> 1800	247.19 ± 145.87		0.3	17.62 ± 1.09	> 1800	> 1800
	0.5	9.61 ± 0.73	> 1800	175.90 ± 24.68		0.5	27.39 ± 3.40	> 1800	> 1800
	0.7	4.88 ± 0.91	138.00 ± 127.46	152.43 ± 52.32		0.7	16.02 ± 4.01	> 1800	> 1800
	0.9	4.70 ± 0.87	2.51 ± 1.21	124.11 ± 45.91		0.9	38.86 ± 2.74	> 1800	> 1800

Table 4 Runtime comparison in instances with *mixed-integer* here-and-now decisions and ellipsoidal uncertainty set. Entries are mean $\pm$ standard deviation in seconds over five independent replications.

n	Algorithm 3	C&CG	ADR
3	0.17 ± 0.02	0.07 ± 0.04	0.23 ± 0.03
10	0.88 ± 0.35	1.46 ± 0.70	41.95 ± 5.16
30	3.05 ± 0.37	1390.24 ± 572.77	> 1800

5.2. Mixed-integer Here-and-Now Decisions

We now test instances with integrality restrictions on the here-and-now decision variables. To this end, we add facility-opening decisions $y_i \in \{0, 1\}$, each incurring a setup cost f_i , to the adaptive robust location-transportation problem. We further impose constraints $\text{capacity}_i \leq K_i y_i$ for all $i \in [m]$, which ensures that capacities can only be installed at opened facilities. We follow the same approach to generate instances as in Section 5.1 and we sample f_i and K_i uniformly from the integer ranges $[300, 400]$ and $[800, 1200]$, respectively.

We report the computing times (all in wall clock seconds) of solving (ARO) instances through Algorithm 3 using $w = 1$ and the uniform proposal distribution. When implementing Algorithm 3, we compute the true objective value of the solution x^* returned by the branch-and-cut method. Across all tested instances, each ensuing objective value is within 0.1% of the true optimal value.

Ellipsoidal uncertainty set. We report the computing times with ellipsoidal uncertainty sets in Table 4, where $n \in \{3, 10, 30\}$. From this table, we observe that Algorithm 3 demonstrates better scalability than C&CG and ADR as n increases. For example, the average runtime of Algorithm 3 becomes at least an order of magnitude shorter than those of C&CG and ADR when n reaches 30. This demonstrates that, although we need to maintain multiple Markov chains in a branch-and-cut tree, the soft separation still manages to identify adversarial scenarios at a relatively lower computational burden.

Table 5 Runtime comparison in instances with *mixed-integer* here-and-now decisions and scenario uncertainty set. Entries are mean $\pm$ standard deviation in seconds over five independent replications.

$ \mathcal{I} $	$n = 3$			$n = 10$			$n = 30$		
	Algorithm 3	C&CG	ADR	Algorithm 3	C&CG	ADR	Algorithm 3	C&CG	ADR
1,000	0.32 $\pm$ 0.12	0.18 $\pm$ 0.08	1.38 $\pm$ 0.55	1.47 $\pm$ 0.45	1.74 $\pm$ 0.71	373.22 $\pm$ 90.40	3.60 $\pm$ 0.25	2.85 $\pm$ 0.48	> 1800
5,000	0.37 $\pm$ 0.15	0.64 $\pm$ 0.31	8.88 $\pm$ 2.37	1.44 $\pm$ 0.48	5.01 $\pm$ 1.54	1397.60 $\pm$ 594.33	3.64 $\pm$ 0.51	12.63 $\pm$ 0.65	> 1800
10,000	0.35 $\pm$ 0.18	1.37 $\pm$ 0.69	25.72 $\pm$ 7.85	1.37 $\pm$ 0.44	11.75 $\pm$ 6.75	> 1800	3.58 $\pm$ 0.38	24.88 $\pm$ 1.41	> 1800
50,000	0.33 $\pm$ 0.15	6.58 $\pm$ 3.44	118.33 $\pm$ 28.03	2.61 $\pm$ 0.60	64.59 $\pm$ 28.64	> 1800	3.84 $\pm$ 0.42	151.00 $\pm$ 38.73	> 1800
100,000	0.42 $\pm$ 0.17	14.49 $\pm$ 6.41	256.19 $\pm$ 99.19	1.98 $\pm$ 0.56	82.62 $\pm$ 17.53	> 1800	3.82 $\pm$ 0.45	238.43 $\pm$ 12.47	> 1800

Table 6 Runtime comparison for instances with mixed-integer here-and-now decisions and budget uncertainty. Times are reported as mean $\pm$ standard deviation in seconds over five independent replications.

n	Γ	Algorithm 3	C&CG	ADR	n	Γ	Algorithm 3	C&CG	ADR
30	0.1	4.73 $\pm$ 1.25	0.28 $\pm$ 0.05	24.84 $\pm$ 12.95	70	0.1	38.97 $\pm$ 6.40	889.39 $\pm$ 313.63	> 1800
	0.3	4.41 $\pm$ 0.84	5.45 $\pm$ 3.19	59.31 $\pm$ 25.07		0.3	40.35 $\pm$ 8.64	> 1800	> 1800
	0.5	3.73 $\pm$ 0.31	5.09 $\pm$ 3.41	92.74 $\pm$ 20.12		0.5	35.39 $\pm$ 2.47	> 1800	> 1800
	0.7	3.86 $\pm$ 0.40	0.93 $\pm$ 1.23	113.54 $\pm$ 38.49		0.7	36.57 $\pm$ 3.00	449.30 $\pm$ 473.53	> 1800
	0.9	4.03 $\pm$ 0.65	0.19 $\pm$ 0.03	104.89 $\pm$ 31.09		0.9	43.07 $\pm$ 9.19	254.14 $\pm$ 237.96	> 1800
50	0.1	24.58 $\pm$ 13.12	17.37 $\pm$ 7.83	1364.19 $\pm$ 631.61	100	0.1	42.96 $\pm$ 5.80	> 1800	> 1800
	0.3	22.72 $\pm$ 15.61	1515.93 $\pm$ 686.48	1469.35 $\pm$ 747.72		0.3	42.93 $\pm$ 3.89	> 1800	> 1800
	0.5	19.10 $\pm$ 4.07	1307.20 $\pm$ 775.33	> 1800		0.5	46.88 $\pm$ 6.73	> 1800	> 1800
	0.7	17.01 $\pm$ 7.43	88.60 $\pm$ 100.68	> 1800		0.7	46.26 $\pm$ 4.87	1469.80 $\pm$ 795.10	> 1800
	0.9	14.76 $\pm$ 2.50	1.32 $\pm$ 0.84	1316.66 $\pm$ 667.30		0.9	45.15 $\pm$ 6.00	1076.80 $\pm$ 548.60	> 1800

Scenario uncertainty set. We report the computing times with scenario uncertainty sets in Table 5, where $n \in \{3, 10, 30\}$ and $|\mathcal{I}|$ varies from 1,000 to 100,000. This table shows that Algorithm 3 scales favorably with both n and $|\mathcal{I}|$. Indeed, for each fixed n , the runtime of Algorithm 3 remains nearly constant as $|\mathcal{I}|$ increases. In addition, when $n = 30$, the ADR runtime exceeds the time limit in all replications, while Algorithm 3 is able to solve all instances within the time limit.

Budget uncertainty set. Finally, we report the computing times with the budget uncertainty set in Table 6, where n varies from 30 to 100 and Γ varies between 0.1 and 0.9. From this table, we observe that C&CG performs the best when n is small and Γ is close to binary (e.g., when $n = 30, \Gamma = 0.1$), and Algorithm 3 outperforms C&CG and ADR as n increases, especially when $n = 100$. In addition, Algorithm 3 shows better scalability with Γ than C&CG and ADR. This confirms our earlier observation from Table 3.

6. Conclusion

This paper studied alternative algorithms for solving adaptive robust optimization (ARO) based on soft, rather than exact, separation. Soft separation adapts a time-inhomogeneous Markov chain to optimization iterates and generates high-confidence adversarial scenarios for ARO. We applied the soft separation to solve (ARO) with both continuous and mixed-integer here-and-now decisions. In the continuous setting, we established polynomial-time convergence of a projected stochastic subgradient descent algorithm driven by soft separation; and in the mixed-integer setting, we

derived a high-probability certificate of global optimality for a branch-and-cut method incorporating soft separation. Numerical results demonstrated the effectiveness and scalability of the proposed methods, particularly as the problem dimension and number of scenarios increase.

References

- An Y, Zeng B (2014) Exploring the modeling capacity of two-stage robust optimization: Variants of robust unit commitment model. *IEEE Transactions on Power Systems* 30(1):109–122.
- Asadi K, Littman ML (2017) An alternative softmax operator for reinforcement learning. Precup D, Teh YW, eds., *Proceedings of the 34th International Conference on Machine Learning*, volume 70 of *Proceedings of Machine Learning Research*, 243–252 (PMLR).
- Aykin T (1996) Optimal shift scheduling with multiple break windows. *Management Science* 42(4):591–602.
- Ben-Tal A, Goryashko A, Guslitzer E, Nemirovski A (2004) Adjustable robust solutions of uncertain linear programs. *Mathematical Programming* 99(2):351–376.
- Ben-Tal A, Nemirovski A (1999) Robust solutions of uncertain linear programs. *Operations Research Letters* 25(1):1–13.
- Ben-Tal A, Nemirovski A (2001) *Lectures on modern convex optimization: analysis, algorithms, and engineering applications* (SIAM).
- Bennett KP, Mangasarian OL (1993) Bilinear separation of two sets in n-space. *Computational Optimization and Applications* 2(3):207–227.
- Bertsimas D, Biddkhor H (2015) On the performance of affine policies for two-stage adaptive optimization: a geometric perspective. *Mathematical Programming* 153(2):577–594.
- Bertsimas D, Den Hertog D, Pauphilet J (2021) Probabilistic guarantees in robust optimization. *SIAM Journal on Optimization* 31(4):2893–2920.
- Bertsimas D, Goyal V (2012) On the power and limitations of affine policies in two-stage adaptive optimization. *Mathematical Programming* 134(2):491–531.
- Bertsimas D, Goyal V (2013) On the approximability of adjustable robust convex optimization under uncertainty. *Mathematical Methods of Operations Research* 77(3):323–343.
- Bertsimas D, Goyal V, Sun XA (2011) A geometric characterization of the power of finite adaptability in multistage stochastic and adaptive optimization. *Mathematics of Operations Research* 36(1):24–54.
- Bertsimas D, Sim M (2004) The price of robustness. *Operations Research* 52(1):35–53.
- Bertsimas D, Sim M, Zhang M (2019) Adaptive distributionally robust optimization. *Management Science* 65(2):604–618.
- Byrne S, Girolami M (2013) Geodesic Monte Carlo on embedded manifolds. *Scandinavian Journal of Statistics* 40(4):825–845.
- Campi MC, Garatti S (2008) The exact feasibility of randomized solutions of uncertain convex programs. *SIAM Journal on Optimization* 19(3):1211–1230.
- Campi MC, Garatti S (2018) Wait-and-judge scenario optimization. *Mathematical Programming* 167(1):155–189.

- Cappé O, Moulines E, Rydén T (2005) *Inference in hidden Markov models* (Springer).
- Che E, Dong J, Tong XT (2026) Stochastic gradient descent with adaptive data. *Operations Research Articles* in Advance.
- Chen X, Sim M, Sun P, Zhang J (2008) A linear decision-based approximation approach to stochastic programming. *Operations Research* 56(2):344–357.
- Cheng C, Qi M, Zhang Y, Rousseau LM (2018) A two-stage robust approach for the reliable logistics network design problem. *Transportation Research Part B: Methodological* 111:185–202.
- Dobrushin RL (1956) Central limit theorem for nonstationary markov chains. i. *Theory of Probability and Its Applications* 1(1):65–80.
- El Housni O, Foussoul A, Goyal V (2024) LP-based approximations for disjoint bilinear and two-stage adjustable robust optimization. *Mathematical Programming* 206(1):239–281.
- El Housni O, Goyal V (2021) On the optimality of affine policies for budgeted uncertainty sets. *Mathematics of Operations Research* 46(2):674–711.
- Gabrel V, Lacroix M, Murat C, Remli N (2014) Robust location transportation problems under uncertain demands. *Discrete Applied Mathematics* 164:100–111.
- Georghiou A, Tsoukalas A, Wiesemann W (2020) A primal–dual lifting scheme for two-stage robust optimization. *Operations Research* 68(2):572–590.
- Georghiou A, Tsoukalas A, Wiesemann W (2026) On the optimality of affine decision rules in distributionally robust optimization. *Management Science* 72(2):1456–1471.
- Han E, Bandi C, Nohadani O (2023) On finite adaptability in two-stage distributionally robust optimization. *Operations Research* 71(6):2307–2327.
- Huang S, Liu K, Zhang ZH (2023) Column-and-constraint-generation-based approach to a robust reverse logistic network design for bike sharing. *Transportation Research Part B: Methodological* 173:90–118.
- Ji M, Wang S, Peng C, Li J (2022) Two-stage robust telemedicine assignment problem with uncertain service duration and no-show behaviours. *Computers & Industrial Engineering* 169:108226.
- Jiang R, Wang J, Guan Y (2011) Robust unit commitment with wind power and pumped storage hydro. *IEEE Transactions on Power Systems* 27(2):800–810.
- Kong Q, Lee CY, Teo CP, Zheng Z (2013) Scheduling arrivals to a stochastic service delivery system using copositive cones. *Operations Research* 61(3):711–726.
- Li Q, Wai HT (2022) State dependent performative prediction with stochastic approximation. *Proceedings of the 25th International Conference on Artificial Intelligence and Statistics*, volume 151, 3164–3186.
- Li X, Liang J, Chen X, Zhang Z (2026) Convergence and inference of stream stochastic gradient descent, with applications to queueing systems and inventory control. *Operations Research Articles* in Advance.

- Lovász L, Vempala SS (2006) Hit-and-run from a corner. *SIAM Journal on Computing* 35(4):985–1005.
- Lu M, Ran L, Shen ZJM (2015) Reliable facility location design under uncertain correlated disruptions. *Manufacturing & Service Operations Management* 17(4):445–455.
- Luan NN, Kim DS, Yen ND (2022) Two optimal value functions in parametric conic linear programming. *Journal of Optimization Theory and Applications* 193(1):574–597.
- Makur A, Singh J (2024) Doeblin coefficients and related measures. *IEEE Transactions on Information Theory* 70(7):4667–4692.
- Minoux M (2011) On 2-stage robust LP with RHS uncertainty: complexity results and applications. *Journal of Global Optimization* 49(3):521–537.
- Moreira A, Piancó F, Fanzeres B, Street A, Jiang R, Zhao C, Heleno M (2024) Distribution system operation amidst wildfire-prone climate conditions under decision-dependent line availability uncertainty. *IEEE Transactions on Power Systems* 39(5):6522–6538.
- Neal RM (2011) MCMC using Hamiltonian dynamics. Brooks S, Gelman A, Jones GL, Meng XL, eds., *Handbook of Markov Chain Monte Carlo*, 113–162 (Chapman and Hall/CRC).
- Qi J (2017) Mitigating delays and unfairness in appointment systems. *Management Science* 63(2):566–583.
- Qi M, Jiang R, Shen S (2024) Sequential competitive facility location: Exact and approximate algorithms. *Operations Research* 72(1):300–316.
- Roberts GO, Tweedie RL (1996) Exponential convergence of langevin distributions and their discrete approximations. *Bernoulli* 2(4):341–363.
- Rudin W (1987) *Real and complex analysis, 3rd ed.* (USA: McGraw-Hill, Inc.), ISBN 0070542341.
- Ryu M, Jiang R (2025) Nurse staffing under absenteeism: A distributionally robust optimization approach. *Manufacturing & Service Operations Management* 27(2):624–639.
- Sutton RS, Barto AG (2018) *Reinforcement Learning: An Introduction* (MIT Press), 2 edition.
- Tsang MY, Shehadeh KS, Curtis FE (2023) An inexact column-and-constraint generation method to solve two-stage robust optimization problems. *Operations Research Letters* 51(1):92–98.
- Yuan W, Wang J, Qiu F, Chen C, Kang C, Zeng B (2016) Robust optimization-based resilient distribution network planning against natural disasters. *IEEE Transactions on Smart Grid* 7(6):2817–2826.
- Zappa E, Holmes-Cerfon M, Goodman J (2018) Monte Carlo on manifolds: Sampling densities and integrating functions. *Communications on Pure and Applied Mathematics* 71(12):2609–2647.
- Zeng B, Wang W (2022) Two-stage robust optimization with decision dependent uncertainty. URL <https://arxiv.org/pdf/2203.16484>.
- Zeng B, Zhao L (2013) Solving two-stage robust optimization problems using a column-and-constraint generation method. *Operations Research Letters* 41(5):457–461.

- Zhao L, Zeng B (2012) An exact algorithm for two-stage robust optimization with mixed integer recourse problems. URL <https://optimization-online.org/wp-content/uploads/2012/01/3310.pdf>.
- Zhen J, Den Hertog D, Sim M (2018) Adjustable robust optimization via Fourier–Motzkin elimination. *Operations Research* 66(4):1086–1100.

Supplementary Results

EC.1. Extension to the General Conic (ARO)

This section extends the analysis in Section 2 to a general cone $\mathcal{K}$ in (ARO). The main difference is that relatively complete recourse alone no longer guarantees strong duality and a uniform bound of the dual optimal solutions. We retain Assumptions 1–2 and replace Assumption 3 with the following conic regularity condition to prove counterparts of Lemmas 1 and 3.

ASSUMPTION EC.1 (Conic recourse regularity). *The set $\mathcal{K} \subseteq \mathbb{R}^m$ is a nonempty closed convex cone with nonempty interior. For every $(x, \xi) \in \mathcal{X} \times \Xi$, there exists a point $\bar{y} = \bar{y}(x, \xi)$ such that*

$$G\bar{y} + Ex + U\xi - h \in \text{int } \mathcal{K}. \quad (\text{EC.1})$$

Define $t(x, \xi) := h - Ex - U\xi$, $\Pi_{\mathcal{K}} := \{\pi \in \mathcal{K}^* : G^\top \pi = b\}$, where $\mathcal{K}^* := \{\pi \in \mathbb{R}^m : \pi^\top z \geq 0 \text{ for all } z \in \mathcal{K}\}$ is the dual cone. By conic strong duality and Assumption 2, (EC.1) implies that

$$Q(x, \xi) = \max_{\pi \in \Pi_{\mathcal{K}}} \pi^\top t(x, \xi), \quad \forall (x, \xi) \in \mathcal{X} \times \Xi, \quad (\text{EC.2})$$

and the maximum is attained (see, e.g., Ben-Tal and Nemirovski 2001). The regularity properties of $Q(x, \xi)$ follow.

LEMMA EC.1 (Conic counterparts of Lemma 1). *Let $\pi^*(x, i)$ denote the minimum-norm maximizer of (EC.2) at $(x, \xi^{(i)})$. Then, $-E^\top \pi^*(x, i) \in \partial_x Q(x, \xi^{(i)})$. Furthermore, there exists a constant $L_Q > 0$ such that, with $L'_Q := L_Q \|E\|$ and $L'_\xi := L_Q \|U\|$, the following hold:*

- (1) *Lipschitz continuity in x : For every $\xi \in \Xi$, $Q(\cdot, \xi)$ is L'_Q -Lipschitz continuous on $\mathcal{X}$;*
- (2) *Lipschitz continuity in ξ : For every $x \in \mathcal{X}$, $Q(x, \cdot)$ is L'_ξ -Lipschitz continuous on Ξ ;*
- (3) *Bounded subgradient selection: The selected dual optimizer is Borel measurable and satisfies $\|E^\top \pi^*(x, i)\| \leq L'_Q$ for every $x \in \mathcal{X}$ and every scenario $\xi^{(i)}$.*

Proof of Lemma EC.1: First, dual feasibility and optimality of $\pi^*(x, i)$ give, for every $x' \in \mathcal{X}$, $Q(x', \xi^{(i)}) \geq \pi^*(x, i)^\top (h - Ex' - U\xi^{(i)}) = Q(x, \xi^{(i)}) + (-E^\top \pi^*(x, i))^\top (x' - x)$, which proves $-E^\top \pi^*(x, i) \in \partial_x Q(x, \xi^{(i)})$. It remains to establish a bound that is uniform over $\mathcal{X} \times \Xi$. Consider the right-hand-side value function $\vartheta(t) := \inf_y \{b^\top y : Gy - t \in \mathcal{K}\}$. For every $t(x, \xi^{(i)})$ generated by $x \in \mathcal{X}$, Assumption EC.1 places t in the interior of the domain on which ϑ is finite. Hence, ϑ is locally Lipschitz at these right-hand sides (see Luan et al. 2022). Because $\mathcal{X}$ is compact and Ξ is finite, the sets $\{t(x, \xi^{(i)}) : x \in \mathcal{X}\}$ with $i \in \mathcal{I}$ form a finite collection of compact sets. Moreover, since the dual feasible set does not depend on the right-hand side, $\pi^*(x, i)$ remains dual feasible at t' . Weak duality and optimality at $t(x, \xi^{(i)})$ therefore give

$$\vartheta(t') \geq \pi^*(x, i)^\top t' = \pi^*(x, i)^\top t(x, \xi^{(i)}) + \pi^*(x, i)^\top (t' - t(x, \xi^{(i)})) = \vartheta(t(x, \xi^{(i)})) + \pi^*(x, i)^\top (t' - t(x, \xi^{(i)})),$$

where the last equality follows from strong duality. It follows that $\pi^*(x, i) \in \partial\vartheta(t(x, \xi^{(i)}))$.

Local boundedness of the subdifferential of a finite convex function therefore gives a common constant $L_Q < \infty$ such that every dual optimizer $\pi^*(x, i)$ in (EC.2) satisfies $\|\pi^*(x, i)\| \leq L_Q$. Setting $L'_Q := L_Q\|E\|$ and $L'_\xi := L_Q\|U\|$ gives $\|E^\top \pi^*(x, i)\| \leq L'_Q$ and the Lipschitz continuity in claims (1)–(2). $\square$

Lemma EC.1 supplies exactly the properties used in the convergence analysis of Section 2. In particular, with $Z_x(i) := g_f(x) - E^\top \pi^*(x, i)$, $z_w(x) := \sum_{i \in \mathcal{I}} p_x(i) Z_x(i) \in \partial F_w(x)$, we have $\|Z_x(i)\| \leq L_f + L'_Q$, $\|z_w(x)\| \leq L_f + L'_Q$. Consequently, Proposition 1 remains valid and the convergence analysis of Algorithm 1 carries over without further changes.

EC.2. Adaptive Proposal for the Budget Uncertainty Set

We detail the exact samplers for the budget uncertainty sets in Section 5.1. For budget Γ , each extreme point of the budget-saturating face (which suffices for our purposes since $Q(x, \xi)$ is componentwise nondecreasing in ξ) has $\lfloor n\Gamma \rfloor$ coordinates equal to one and one coordinate equal to $\Gamma - \lfloor n\Gamma \rfloor \in [0, 1)$. Define

$$r_\Gamma^2 = \lfloor n\Gamma \rfloor + (n\Gamma - \lfloor n\Gamma \rfloor)^2 - \frac{(n\Gamma)^2}{n}, \quad \text{index}_x(i) = \frac{\kappa (\text{Id} - \frac{1}{n} \mathbf{1}\mathbf{1}^\top) (-U^\top \pi^*(x, i))}{r_\Gamma \| (\text{Id} - \frac{1}{n} \mathbf{1}\mathbf{1}^\top) (-U^\top \pi^*(x, i)) \|}.$$

The proposal distribution over the extreme points is $q_x^{\text{budget}}(\xi' | \xi^{(i)}) \propto \exp(\text{index}_x(i)^\top \xi')$. We use a dynamic program (DP) to compute the denominator (normalizer) of $q_x^{\text{budget}}(\xi' | \xi^{(i)})$ without enumerating the extreme points. We let $Z_j(r, h)$ be the total weight obtainable from coordinates $j, \dots, n$ when r unit coordinates and $h \in \{0, 1\}$ fractional coordinates remain. With $Z_{n+1}(0, 0) = 1$ and all other terminal states equal to zero, the backward recursion of the DP is defined through

$$Z_j(r, h) = Z_{j+1}(r, h) + \mathbf{1}_{\{r>0\}} e^{\text{index}_x(j)} Z_{j+1}(r-1, h) + \mathbf{1}_{\{h=1\}} e^{(n\Gamma - \lfloor n\Gamma \rfloor) \text{index}_x(j)} Z_{j+1}(r, 0)$$

Then, $\sum_{i \in \mathcal{I}} \exp(\text{index}_x(i)^\top \xi') = Z_1(\lfloor n\Gamma \rfloor, \mathbf{1}_{\{n\Gamma - \lfloor n\Gamma \rfloor > 0\}})$. Thus, a sample is generated sequentially by assigning coordinate j the value 0, 1, or $n\Gamma - \lfloor n\Gamma \rfloor > 0$ with weight $Z_{j+1}(r, h)/Z_j(r, h)$, $(\mathbf{1}_{\{r>0\}} e^{\text{index}_x(j)} Z_{j+1}(r-1, h))/Z_j(r, h)$, or $(\mathbf{1}_{\{h=1\}} e^{(n\Gamma - \lfloor n\Gamma \rfloor) \text{index}_x(j)} Z_{j+1}(r, 0))/Z_j(r, h)$ respectively. After the draw, the DP state is updated through

$$(r, h) \leftarrow \begin{cases} (r, h), & \text{if } \xi_j = 0, \\ (r-1, h), & \text{if } \xi_j = 1, \\ (r, 0), & \text{if } \xi_j = n\Gamma - \lfloor n\Gamma \rfloor. \end{cases}$$

This DP defines a valid sampling procedure. Indeed, any infeasible assignment has zero probability because the corresponding DP value function equals zero. Finally, we note that building the DP table requires $\mathcal{O}(n\lfloor n\Gamma \rfloor)$ operations and the subsequent draw requires $\mathcal{O}(n)$ operations.

EC.3. Supplementary Implementation Details

1. *Preconditioning subgradients.* When implementing Algorithm 1 in Section 5.1, we precondition the subgradients $Z_{x_n}(I_{n+1}) = g_f(x) - E^\top \pi^*(x, \xi^{I_{n+1}})$ to reduce the sensitivity to heterogeneous magnitudes of the subgradient entries. Recall that, in the robust location-transportation problem, the first-stage cost is $f(x) := c^\top x$ and the coefficients $c \in \mathbb{R}^{2m}$ consist of two parts, with the first m entries denoting the facility opening costs and the last m entries denoting the capacity expansion costs. Then, $g_f(x) \equiv c$. In addition, the first m entries of the dual vector $E^\top \pi^*(x, \xi^{I_{n+1}})$ equal zero because the facility opening variables y_i do not partake the second-stage transportation linear program, and the last m entries of $E^\top \pi^*(x, \xi^{I_{n+1}})$ lie within the interval $[0, h\mathbf{1}]$, where h denotes the unit cost of using the emergency supply. Therefore, we scale $Z_{x_n}(I_{n+1})$ by a diagonal matrix $H \in \mathbb{R}^{2m \times 2m}$ with

$$H_{ii} \propto \max\{|c_i|, 10^{-8}\}, \quad H_{m+i, m+i} \propto \max\{|c_{m+i}|, |c_{m+i} - h|, 10^{-8}\} \quad \forall i \in [m],$$

and we normalize H so that $\det(H) = 1$. Accordingly, we implement the subgradient descent as

$$x \leftarrow \Pi_{\mathcal{X}}(x - \alpha_{n+1} H^{-1} Z_{x_n}(I_{n+1})),$$

where we pre-compute H and fix it throughout all experiments.

2. *Practical certification.* When implementing Algorithm 3 in Section 5.2, we generalize the certification criterion $\bar{\theta} \geq Q^*(\bar{x})$ to be $\bar{\theta} + \varepsilon_{\text{cert}} \geq Q^*(\bar{x})$ for a certification error $\varepsilon_{\text{cert}} \in (0, \varepsilon_{\text{opt}})$. Accordingly, if we establish a slightly more conservative upper bound $\hat{U} := f(\bar{x}) + \bar{\theta} + \varepsilon_{\text{cert}}$, then the same probabilistic guarantee of Theorem 5 continues to hold. Practically, the looser certification criterion makes it easier to find an adversarial scenario, giving rise to a smaller value of ρ . For our independent uniform proposal ν , the corresponding one-call miss probability becomes

$$\rho_{\varepsilon_{\text{cert}}}(\bar{x}, \bar{\theta}) := 1 - \nu(\{\xi \in \Xi : Q(\bar{x}, \xi) > \bar{\theta} + \varepsilon_{\text{cert}}\}).$$

In practice, $\rho_{\varepsilon_{\text{cert}}}(\bar{x}, \bar{\theta})$ can be estimated from the empirical quantiles of the recourse function values $\{Q(x_n, \xi^{(I_{n+1})})\}_n$, which are observed in the warm-start of Algorithm 3 (line 1). We use $\rho = 0.9$ and $\varepsilon_{\text{cert}} = 100$ for all experiments in Section 5.2.

EC.4. Technical Proofs

EC.4.1. Proof of Lemma 1

Proof: First, to show the representation (2), we write the dual formulation for $Q(x, \xi)$ as $\max_{\pi \in \Pi} \pi^\top (h - Ex - U\xi)$, where π denotes dual variables associated with the constraints in (1) and $\Pi = \{\pi \geq 0 : G^\top \pi = b\}$ is the dual feasible region. The strong duality holds because of Assumption 2.

Since Π is a nonempty pointed polyhedron, its set of extreme points $\mathcal{V} := \text{ext}(\Pi)$ is nonempty and finite. It follows from Assumption 2 that, for each $Q(x, \xi)$, an optimal π^* exists in $\mathcal{V}$.

For *fixed* (x, i) , the auxiliary π^* -selection problem admits a linear programming formulation $\min_{\pi \in \Pi, z} \{ \mathbf{1}^\top z : \pi^\top (h - Ex - U\xi) = Q(x, \xi^{(i)}), -z \leq E^\top \pi \leq z \}$. This program is feasible and bounded below by zero, so its minimum is attained. More generally, for *every* nonempty face (denoted by **Face**) of Π , the same argument shows that $\min_{\pi \in \text{Face}} \|E^\top \pi\|_1$ is attained. We fix, for each **Face**, a representative $\hat{\pi}_{\text{Face}} \in \arg \min_{\pi \in \text{Face}} \|E^\top \pi\|_1$. The dual optimal set is a nonempty face of Π , denoted by **Opt-Face**, so the prescribed selection is $\pi^*(x, i) = \hat{\pi}_{\text{Opt-Face}}$. **Note that this mapping is measurable.** To see this, Π has finitely many vertices and extreme rays. The region of (x, ξ) on which **Opt-Face** equals a given face is determined by finitely many equalities and strict inequalities involving their affine objective values, and is therefore Borel measurable, denoted by R_{Face} . On each such Borel region, the selection equals the fixed representative $\hat{\pi}_{\text{Face}}$ of the corresponding face. Hence, the preimage of any open set under π^* , namely $(\pi^*)^{-1}(\mathcal{O})$, is a finite union of such Borel regions (i.e. $\cup_{\text{Face}: \hat{\pi}_{\text{Face}} \in \mathcal{O}} R_{\text{Face}}$), which is a Borel set and proves measurability.

Second, for every $x, y \in \mathcal{X}$ and every scenario $\xi^{(i)}$, dual feasibility and optimality at x give

$$\begin{aligned} Q(y, \xi^{(i)}) &= \max_{\pi \in \Pi} \pi^\top (h - Ey - U\xi^{(i)}) \\ &\geq \pi^*(x, \xi^{(i)})^\top (h - Ey - U\xi^{(i)}) \\ &= Q(x, \xi^{(i)}) + (-E^\top \pi^*(x, \xi^{(i)}))^\top (y - x). \end{aligned}$$

Thus, $-E^\top \pi^*(x, i) \in \partial_x Q(x, \xi^{(i)})$. It remains to show properties (1)–(3). Choose $L_Q := \max \{1, \max_{v \in \mathcal{V}} \|v\|, \max \|\hat{\pi}_{\text{Face}}\|\} < \infty$. Then, for every $x_1, x_2 \in \mathcal{X}$ and $\xi \in \Xi$, the finite maximum representation yields $|Q(x_1, \xi) - Q(x_2, \xi)| \leq \max_{v \in \mathcal{V}} |v^\top E(x_1 - x_2)| \leq L_Q \|E\| \|x_1 - x_2\|$, proving that $Q(\cdot, \xi)$ is L'_Q -Lipschitz on $\mathcal{X}$. Similarly, for all $\xi^1, \xi^2 \in \Xi$, $|Q(x, \xi^1) - Q(x, \xi^2)| \leq \max_{v \in \mathcal{V}} |v^\top U(\xi^1 - \xi^2)| \leq L_Q \|U\| \|\xi^1 - \xi^2\|$, proving that $Q(x, \cdot)$ is L'_ξ -Lipschitz on Ξ .

Finally, since $\pi^*(x, i)$ equals one of the fixed representatives $\hat{\pi}_{\text{Face}}$, it holds that $\|E^\top \pi^*(x, i)\| \leq \|E\| \|\pi^*(x, i)\| \leq L_Q \|E\| = L'_Q$. Together with measurability, this proves property (3). $\square$

EC.4.2. Proof of Lemma 3

Proof For any $x, y \in \mathcal{X}$, the definition of p_x gives

$$\begin{aligned} F_w(y) - F_w(x) &= f(y) - f(x) + w \log \left(\sum_{i \in \mathcal{I}} p_x(i) \exp \left(\frac{Q(y, \xi^{(i)}) - Q(x, \xi^{(i)})}{w} \right) \right) \\ &\geq f(y) - f(x) + \sum_{i \in \mathcal{I}} p_x(i) (Q(y, \xi^{(i)}) - Q(x, \xi^{(i)})) \\ &\geq g_f(x)^\top (y - x) + \sum_{i \in \mathcal{I}} p_x(i) (-E^\top \pi^*(x, i))^\top (y - x) = z_w(x)^\top (y - x), \end{aligned}$$

where the first inequality follows from Jensen's inequality and the second follows from the subgradient inequalities for f and $Q(\cdot, \xi^{(i)})$ by Lemma 1. Thus, $z_w(x) \in \partial F_w(x)$.

For any $x \in \mathcal{X}$ and any unit vector u , convexity and Lipschitz continuity on an open neighborhood of $\mathcal{X}$ give $t \cdot g_f(x)^\top u \leq f(x + tu) - f(x) \leq L_f t$ for a sufficiently small $t > 0$. Thus, $\|g_f(x)\| \leq L_f$ and $\|Z_x(i)\| = \|g_f(x) - E^\top \pi^*(x, i)\| \leq \|g_f(x)\| + \|E^\top \pi^*(x, i)\| \leq L_f + L'_Q$. Since $p_x(i) \geq 0$ and $\sum_{i \in \mathcal{I}} p_x(i) = 1$, we have $\|z_w(x)\| = \|\sum_{i \in \mathcal{I}} p_x(i) Z_x(i)\| \leq \sum_{i \in \mathcal{I}} p_x(i) \|Z_x(i)\| \leq L_f + L'_Q$. This completes the proof. $\square$

EC.4.3. Proof of Lemma 4

Proof: For any stochastic matrix S , the coefficient of ergodicity (Dobrushin coefficient, see [Dobrushin \(1956\)](#)) is defined by

$$\tau(S) := \frac{1}{2} \max_{i,k} \|S(i, \cdot) - S(k, \cdot)\|_1 = \max_{\substack{\tilde{z}^\top \mathbf{1} = 0 \\ \|\tilde{z}\|_1 = 1}} \|S^\top \tilde{z}\|_1.$$

Define $R_x := (\text{Id} - P_x + \mathbf{1}p_x^\top)^{-1}$. Since P_x is an irreducible and aperiodic Markov matrix on a finite state space, it has a simple eigenvalue 1 with right eigenvector $\mathbf{1}$ and left eigenvector $p_x^\top$, while all other eigenvalues λ satisfy $|\lambda| < 1$.

For any eigenvector v with $P_x v = \lambda v$ and $\lambda \neq 1$, multiplying by $p_x^\top$ gives $\lambda p_x^\top v = p_x^\top P_x v = p_x^\top v$, hence $(\lambda - 1)p_x^\top v = 0$, and therefore $p_x^\top v = 0$. Moreover, $\mathbb{R}^n = \text{span}\{\mathbf{1}\} \oplus \ker(p_x^\top)$, and both subspaces are invariant under P_x . On $\text{span}\{\mathbf{1}\}$, the operator $P_x - \mathbf{1}p_x^\top$ acts as zero, while on $\ker(p_x^\top)$ it agrees with P_x . Consequently, $\sigma(P_x - \mathbf{1}p_x^\top) = \{0\} \cup (\sigma(P_x) \setminus \{1\})$ and hence $\rho(P_x - \mathbf{1}p_x^\top) = \max\{|\lambda| : \lambda \in \sigma(P_x) \setminus \{1\}\} < 1$. Now fix $N \geq 0$. For the finite sum, we have the exact geometric identity $(\text{Id} - P_x + \mathbf{1}p_x^\top) \sum_{t=0}^N (P_x - \mathbf{1}p_x^\top)^t = \text{Id} - (P_x - \mathbf{1}p_x^\top)^{N+1}$ and similarly on the right.

Since $\rho(P_x - \mathbf{1}p_x^\top) < 1$, we have $\|(P_x - \mathbf{1}p_x^\top)^{N+1}\| \rightarrow 0$ as $N \rightarrow \infty$ in any operator norm. Taking limits yields $R_x = (\text{Id} - P_x + \mathbf{1}p_x^\top)^{-1} = \sum_{t=0}^{\infty} (P_x - \mathbf{1}p_x^\top)^t$. We next bound the $\infty \rightarrow \infty$ norm of R_x . For any vector $z \in \mathbb{R}^n$, define $\tilde{z} := z - (z^\top \mathbf{1})p_x$, then $\tilde{z}^\top \mathbf{1} = 0$. Moreover, $z^\top (P_x - \mathbf{1}p_x^\top) = \tilde{z}^\top P_x$, and, by induction, for every $t \geq 1$, $z^\top (P_x - \mathbf{1}p_x^\top)^t = \tilde{z}^\top P_x^t$. Since p_x is a probability vector, it holds that $\|\tilde{z}\|_1 = \|z - (z^\top \mathbf{1})p_x\|_1 \leq \|z\|_1 + |z^\top \mathbf{1}| \|p_x\|_1 \leq 2\|z\|_1$. Thus, $\|z\|_1 = 1$ implies $\|\tilde{z}\|_1 \leq 2$. Using the row-vector characterization of the induced $\infty \rightarrow \infty$ norm, for every $t \geq 1$,

$$\begin{aligned} \|(P_x - \mathbf{1}p_x^\top)^t\|_{\infty \rightarrow \infty} &= \sup_{\|z\|_1=1} \|z^\top (P_x - \mathbf{1}p_x^\top)^t\|_1 \\ &= \sup_{\|z\|_1=1} \|\tilde{z}^\top P_x^t\|_1 \leq \sup_{\substack{\tilde{z}^\top \mathbf{1}=0 \\ \|\tilde{z}\|_1 \leq 2}} \|\tilde{z}^\top P_x^t\|_1 = 2\tau(P_x^t). \end{aligned}$$

By the submultiplicativity of the Dobrushin coefficient, $\tau(P_x^t) \leq [\tau(P_x)]^t$ for all $t \geq 1$. Therefore, $\|(P_x - \mathbf{1}p_x^\top)^t\|_{\infty \rightarrow \infty} \leq 2[\tau(P_x)]^t$. Finally, using the geometric representation of R_x and the triangle inequality, we conclude that

$$\begin{aligned} \|R_x\|_{\infty \rightarrow \infty} &\leq 1 + \sum_{t=1}^{\infty} \|(P_x - \mathbf{1}p_x^\top)^t\|_{\infty \rightarrow \infty} \\ &\leq 1 + 2 \sum_{t=1}^{\infty} [\tau(P_x)]^t = \frac{1 + \tau(P_x)}{1 - \tau(P_x)} \leq \frac{2}{1 - \tau(P_x)}. \end{aligned}$$

We further bound it through a lemma.

LEMMA EC.2 (Doeblin minorization implies bounded coefficient). *If P_x satisfies the Doeblin minorization condition $P_x(i, \cdot) \geq \gamma_x \nu_x(\cdot)$, for some probability measure ν_x and constant $\gamma_x > 0$, then the ergodicity coefficient satisfies $\tau(P_x) \leq 1 - \gamma_x$.*

This is the standard result from literature; see (Makur and Singh 2024, Eq. (2) and Theorem 1, Part 6). Combining all completes the proof. $\square$

EC.4.4. Proof of Proposition 1

The proof of Proposition 1 uses the following three lemmas. Lemmas EC.3–EC.4 prove Lipschitz continuity of the Markov kernel $P_x(\cdot, \cdot)$, and then Lemma EC.7 bounds the sample error. In addition, Lemma EC.5 generalizes Lemmas EC.3–EC.4 to a general state space. Lemma EC.5 is not directly used in the proof of Proposition 1 but will be useful later in Section 3. Lemma EC.6 establishes the Lipschitz continuity of p_x on a general state space and is used both below and in Section 3.

LEMMA EC.3 (Element/row-wise Lipschitz continuity of the Barker kernel). *Consider a finite state space $\mathcal{I}$ and a family of target distributions $p_x(i) \propto \exp(Q_i(x)/w)$, where each function $Q_i(\cdot)$ is L'_Q -Lipschitz on $\mathcal{X}$ and $w > 0$ is fixed. Let $q(j|i) > 0$ denote a fixed proposal kernel, and define the Barker transition matrix $P_x = (P_x(i, j))_{i, j \in \mathcal{I}}$ by*

$$P_x(i, j) = \begin{cases} q(j|i) A_x(i, j), & j \neq i, \\ 1 - \sum_{k \neq i} q(k|i) A_x(i, k), & j = i, \end{cases} \quad \text{where} \quad A_x(i, j) = \frac{r_x(i, j)}{1 + r_x(i, j)},$$

and $r_x(i, j) := \frac{p_x(j) q(i|j)}{p_x(i) q(j|i)}$. Then for all $x, x' \in \mathcal{X}$ and $i, j \in \mathcal{I}$, $|P_x(i, j) - P_{x'}(i, j)| \leq (L'_Q/(2w)) \|x - x'\|$. Moreover, for each row i , $\sum_{j \in \mathcal{I}} |P_x(i, j) - P_{x'}(i, j)| \leq L_P \|x - x'\|$ with $L_P := L'_Q/w$.

Proof of Lemma EC.3: Write $A_x(i, j) = \sigma(\alpha_x(i, j))$ with $\sigma(t) = \frac{e^t}{1+e^t}$ and $\alpha_x(i, j) = \log p_x(j) - \log p_x(i) + \log q(i|j) - \log q(j|i) = (Q_j(x) - Q_i(x))/w + \log q(i|j) - \log q(j|i)$. Since the proposal is independent of x and each Q_i is L'_Q -Lipschitz, we have

$$|\alpha_x(i, j) - \alpha_{x'}(i, j)| \leq \frac{|Q_j(x) - Q_j(x')| + |Q_i(x) - Q_i(x')|}{w} \leq \frac{2L'_Q}{w} \|x - x'\|.$$

Since $\sigma'(t) = \sigma(t)(1 - \sigma(t)) \leq 1/4$, we have $|A_x(i, j) - A_{x'}(i, j)| \leq \frac{1}{4} |\alpha_x(i, j) - \alpha_{x'}(i, j)| \leq (L'_Q/(2w)) \|x - x'\|$. We note that, for off-diagonal entries,

$$|P_x(i, j) - P_{x'}(i, j)| = q(j | i) |A_x(i, j) - A_{x'}(i, j)| \leq q(j | i) \frac{L'_Q}{2w} \|x - x'\| \leq \frac{L'_Q}{2w} \|x - x'\|,$$

and for the diagonal entry,

$$|P_x(i, i) - P_{x'}(i, i)| \leq \sum_{k \neq i} q(k | i) |A_x(i, k) - A_{x'}(i, k)| \leq \frac{L'_Q}{2w} \|x - x'\| \sum_{k \neq i} q(k | i) \leq \frac{L'_Q}{2w} \|x - x'\|.$$

Summing over j gives $\sum_j |P_x(i, j) - P_{x'}(i, j)| \leq \frac{L'_Q}{2w} \|x - x'\| (\sum_{k \neq i} q(k | i) + 1) \leq \frac{L'_Q}{w} \|x - x'\|$. $\square$

LEMMA EC.4 (Element/row-wise Lipschitz continuity of the MH kernel). *Let $\mathcal{I}$ be a finite state space and consider target distributions $p_x(i) \propto \exp(Q_i(x)/w)$, where each $Q_i(\cdot)$ is L'_Q -Lipschitz on $\mathcal{X}$ and $w > 0$ is fixed. Let $q(j | i) > 0$ be a proposal kernel with $\sum_j q(j | i) = 1$ for all i . Define the Metropolis–Hastings transition matrix by*

$$P_x(i, j) = \begin{cases} q(j | i) A_x(i, j), & j \neq i, \\ 1 - \sum_{k \neq i} q(k | i) A_x(i, k), & j = i, \end{cases} \quad A_x(i, j) = \min \left\{ 1, \frac{p_x(j) q(i | j)}{p_x(i) q(j | i)} \right\}.$$

Then for all $x, x' \in \mathbb{R}^d$ and $i, j \in \mathcal{I}$, $|P_x(i, j) - P_{x'}(i, j)| \leq q(j | i) (2L'_Q/w) \|x - x'\|$ whenever $j \neq i$ and $|P_x(i, i) - P_{x'}(i, i)| \leq (4L'_Q/w) \|x - x'\|$. In particular, for each row i , it holds that $\sum_{j \in \mathcal{I}} |P_x(i, j) - P_{x'}(i, j)| \leq L_P \|x - x'\|$ with $L_P := 4L'_Q/w$.

Proof of Lemma EC.4: Set $\alpha_x(i, j) := \log p_x(j) - \log p_x(i) + \log q(i | j) - \log q(j | i)$. Note $A_x(i, j) = g(\alpha_x(i, j))$ with $g(t) := \min\{1, e^t\}$. Since g is 1-Lipschitz on $\mathbb{R}$, we have $|A_x(i, j) - A_{x'}(i, j)| \leq |\alpha_x(i, j) - \alpha_{x'}(i, j)|$. Since the proposal is independent of x and $\log p_x(j) - \log p_x(i) = (Q_j(x) - Q_i(x))/w$, Lipschitz continuity of Q_i implies

$$|\alpha_x(i, j) - \alpha_{x'}(i, j)| \leq \frac{|Q_j(x) - Q_j(x')| + |Q_i(x) - Q_i(x')|}{w} \leq \frac{2L'_Q}{w} \|x - x'\|.$$

Therefore, for $j \neq i$, we have $|P_x(i, j) - P_{x'}(i, j)| = q(j | i) |A_x(i, j) - A_{x'}(i, j)| \leq q(j | i) \frac{2L'_Q}{w} \|x - x'\|$, and, for the diagonal entry,

$$|P_x(i, i) - P_{x'}(i, i)| = \left| \sum_{k \neq i} q(k | i) (A_{x'}(i, k) - A_x(i, k)) \right| \leq \sum_{k \neq i} q(k | i) |A_x(i, k) - A_{x'}(i, k)| \leq \frac{2L'_Q}{w} \|x - x'\|.$$

Summing over j in row i and using $\sum_{j \neq i} q(j | i) \leq 1$ yield the row-wise bound. $\square$

LEMMA EC.5 (Lipschitz continuous kernels for a general state space). *Under Assumption 5, consider a Markov kernel P_x induced by (i) a proposal distribution $\nu(\cdot)$ on $\mathcal{I}$ (which is independent of x) and (ii) either MH or Barker acceptance. Then, Assumption 6 holds with $L_P = 2L'_Q/w$. In other words, $\sup_{I \in \mathcal{I}} \|P_x(I, \cdot) - P_y(I, \cdot)\|_{\text{TV}} \leq (2L'_Q/w) \|x - y\|$ for all $x, y \in \mathcal{X}$.*

Proof of Lemma EC.5: Let $A_x(I, J)$ denote either MH or Barker acceptance distribution. For every measurable $S \subseteq \mathcal{I}$, we have

$$\begin{aligned} P_x(I, S) - P_y(I, S) &= \int_{\mathcal{I}} A_x(I, J) \mathbf{1}_S(J) \nu(dJ) + \mathbf{1}_S(I) \int_{\mathcal{I}} (1 - A_x(I, J)) \nu(dJ) \\ &\quad - \int_{\mathcal{I}} A_y(I, J) \mathbf{1}_S(J) \nu(dJ) - \mathbf{1}_S(I) \int_{\mathcal{I}} (1 - A_y(I, J)) \nu(dJ) \\ &= \int_{\mathcal{I}} (A_x(I, J) - A_y(I, J)) \mathbf{1}_S(J) \nu(dJ) - \mathbf{1}_S(I) \int_{\mathcal{I}} (A_x(I, J) - A_y(I, J)) \nu(dJ) \\ &= \int_{\mathcal{I}} (\mathbf{1}_S(J) - \mathbf{1}_S(I)) (A_x(I, J) - A_y(I, J)) \nu(dJ). \end{aligned}$$

On the one hand, if $A_x(I, J)$ is the MH acceptance distribution, then $A_x(I, J) = 1 \wedge \exp\left(\frac{\widehat{Q}_x(J) - \widehat{Q}_x(I)}{w}\right)$. Because $u \mapsto 1 \wedge e^u$ is 1-Lipschitz, Assumption 5 gives $|A_x(I, J) - A_y(I, J)| \leq (|\widehat{Q}_x(J) - \widehat{Q}_y(J)| + |\widehat{Q}_x(I) - \widehat{Q}_y(I)|)/w \leq (2L'_Q/w)\|x - y\|$. On the other hand, if $A_x(I, J)$ is the Barker acceptance distribution, then $A_x(I, J) = \exp((\widehat{Q}_x(J) - \widehat{Q}_x(I))/w)/(1 + \exp((\widehat{Q}_x(J) - \widehat{Q}_x(I))/w))$. Because $u \mapsto e^u/(1 + e^u)$ is 1-Lipschitz (in fact, it is $(1/4)$ -Lipschitz), the same conclusion holds that $|A_x(I, J) - A_y(I, J)| \leq (2L'_Q/w)\|x - y\|$. It follows that

$$|P_x(I, S) - P_y(I, S)| \leq \max |\mathbf{1}_S(J) - \mathbf{1}_S(I)| \int_{\mathcal{I}} |A_x(I, J) - A_y(I, J)| \nu(dJ) \leq \frac{2L'_Q}{w} \|x - y\|.$$

Taking the supremum over measurable S and $I \in \mathcal{I}$ completes the proof. $\square$

LEMMA EC.6 (Lipschitz continuity of the target distribution). *Under Assumption 5, the target distribution p_x is Lipschitz continuous in x . Specifically, with $L_p := 2L'_Q/w$, it holds that*

$$\left\| \frac{dp_x}{d\nu} - \frac{dp_y}{d\nu} \right\|_{L^1(\nu)} = 2\|p_x - p_y\|_{\text{TV}} = \|\mathbf{1}p_x - \mathbf{1}p_y\|_{\infty \rightarrow \infty} \leq L_p \|x - y\|$$

for all $x, y \in \mathcal{X}$. In particular, when $\mathcal{I}$ is finite, $\|p_x - p_y\|_1 = \|\mathbf{1}(p_x - p_y)\|_{\infty \rightarrow \infty} \leq L_p \|x - y\|$.

Proof of Lemma EC.6 Fix $x, y \in \mathcal{X}$. For $t \in [0, 1]$, define

$$r_t(I) := \frac{\exp\left(\left[(1-t)Q_\tau(y, \xi_I) + tQ_\tau(x, \xi_I)\right]/w\right)}{\int_{\mathcal{I}} \exp\left(\left[(1-t)Q_\tau(y, \xi_J) + tQ_\tau(x, \xi_J)\right]/w\right) \nu(dJ)}.$$

Differentiation under the integral and Assumption 5 yield

$$\frac{d}{dt} r_t(I) = \frac{r_t(I)}{w} \left(Q_\tau(x, \xi_I) - Q_\tau(y, \xi_I) - \int_{\mathcal{I}} [Q_\tau(x, \xi_J) - Q_\tau(y, \xi_J)] r_t(J) \nu(dJ) \right), \left\| \frac{d}{dt} r_t \right\|_{L^1(\nu)} \leq \frac{2L'_Q}{w} \|x - y\|.$$

Therefore,

$$\left\| \frac{dp_x}{d\nu} - \frac{dp_y}{d\nu} \right\|_{L^1(\nu)} = \|r_1 - r_0\|_{L^1(\nu)} \leq \int_0^1 \left\| \frac{d}{dt} r_t \right\|_{L^1(\nu)} dt \leq \frac{2L'_Q}{w} \|x - y\|.$$

The remaining equalities follow from the definition of total variation and the induced ℓ_∞ -operator norm. The finite-state result follows by identifying the densities with their probability vectors.

LEMMA EC.7. Let $\widehat{Z}_x$ denote the centered solution of the Poisson equation $\widehat{Z}_x - P_x \widehat{Z}_x = e_x$ with $p_x^\top \widehat{Z}_x = 0$. Then, there exists a constant B_Z such that

$$\sup_{x \in \mathcal{X}} \max \{ \|\widehat{Z}_x\|_\infty, \|P_x \widehat{Z}_x\|_\infty \} \leq B_Z,$$

where $B_Z = 2C_* DL'_Q$. Here, $D := \max_{x, y \in \mathcal{X}} \|x - y\|_1$ is the diameter of $\mathcal{X}$, L'_Q is the Lipschitz modulus of $Q(\cdot, \xi)$ in Lemma 1, and C_* is the fundamental matrix norm bound from Assumption 4. Moreover, there exists a constant $L_Z < \infty$ such that $\|\widehat{Z}_x - \widehat{Z}_{x'}\|_\infty \leq L_Z \|x - x'\|$ and $\|P_x \widehat{Z}_x - P_{x'} \widehat{Z}_{x'}\|_\infty \leq (L_P B_Z + L_Z) \|x - x'\|$ for all $x, x' \in \mathcal{X}$, where $L_Z := C_*(2L'_Q + DL'_Q L_P) + 2C_*^2 DL'_Q (L_P + L_p)$. Here, $L_p := 2L'_Q/w$ and L_P comes from either Lemma EC.3 or Lemma EC.4.

Proof of Lemma EC.7: (I) Bounding $\max \{ \|\widehat{Z}_x\|_\infty, \|P_x \widehat{Z}_x\|_\infty \}$. Recall that, by definition, $e_x(i) = Q(x, \xi^{(i)}) - Q(x^*, \xi^{(i)}) - \sum_{j \in \mathcal{I}} p_x(j) (Q(x, \xi^{(j)}) - Q(x^*, \xi^{(j)}))$. By Lipschitz continuity of $Q(\cdot, \xi)$ and the definition of D , we have $|Q(x, \xi^{(i)}) - Q(x^*, \xi^{(i)})| \leq L'_Q \|x - x^*\| \leq DL'_Q$. Hence, $\|e_x\|_\infty \leq 2DL'_Q$. Since $p_x^\top e_x = 0$, the centered Poisson solution is, for every $i \in \mathcal{I}$,

$$\widehat{Z}_x(i) = \sum_{j \in \mathcal{I}} [(\text{Id} - P_x + \mathbf{1} p_x^\top)^{-1}]_{ij} e_x(j), \quad \widehat{Z}_x(i) - \sum_{j \in \mathcal{I}} P_x(i, j) \widehat{Z}_x(j) = e_x(i).$$

Then, by Assumption 4 and for all $i \in \mathcal{I}$, $\|\widehat{Z}_x(i)\|_\infty \leq C_* \sup_{j \in \mathcal{I}} |e_x(j)| \leq 2C_* DL'_Q$. Now, because P_x is a stochastic row matrix, we have $\|(P_x \widehat{Z}_x)(i)\|_\infty = |\sum_{j \in \mathcal{I}} P_x(i, j) \widehat{Z}_x(j)| \leq \sup_{i \in \mathcal{I}} \|\widehat{Z}_x(i)\|_\infty \leq 2C_* DL'_Q$. Setting $B_Z := 2C_* DL'_Q$ completes this part of the proof.

(II) Lipschitz continuity of $\widehat{Z}_x$ and $P_x \widehat{Z}_x$ on $\mathcal{X}$. Throughout this part, we use the infinity vector norm and its induced operator norm for matrices. To proceed, we recall that

- (A) By Lemma 1, $\sup_{i \in \mathcal{I}} |Q(x, \xi^{(i)}) - Q(x', \xi^{(i)})| \leq L'_Q \|x - x'\|$. Moreover, since $x^* \in \mathcal{X}$, we have $\sup_{i \in \mathcal{I}} |Q(x, \xi^{(i)}) - Q(x^*, \xi^{(i)})| \leq DL'_Q$.
- (B) The Poisson operator is uniformly invertible: $\sup_x \|(I - P_x + \mathbf{1} p_x^\top)^{-1}\|_{\infty \rightarrow \infty} \leq C_* < \infty$ by Assumption 4.
- (C) *Row-wise Lipschitz continuity of P_x .* By Lemma EC.3 (Barker) or Lemma EC.4 (MH), for every $i \in \mathcal{I}$, it holds that $\sum_{j \in \mathcal{I}} |P_x(i, j) - P_{x'}(i, j)| \leq L_P \|x - x'\|$ with $L_P = L'_Q/w$ for Barker (resp. $L_P = 2L'_Q/w$ for MH). Equivalently, in the $\infty \rightarrow \infty$ operator norm, $\|P_x - P_{x'}\|_{\infty \rightarrow \infty} \leq L_P \|x - x'\|$.
- (D) *Lipschitz continuity of p_x in ℓ_1 .* By Lemma EC.6, $\|p_x - p_{x'}\|_1 \leq L_p \|x - x'\|$ with $L_p := 2L'_Q/w$. Moreover, $\|\mathbf{1}(p_x - p_{x'})^\top\|_{\infty \rightarrow \infty} = \|p_x - p_{x'}\|_1 \leq L_p \|x - x'\|$.

We first bound $\sup_{i \in \mathcal{I}} |e_x(i) - e_{x'}(i)|$. By the definition of e_x and properties (A), (D) above, we have

$$\begin{aligned} \sup_{i \in \mathcal{I}} |e_x(i) - e_{x'}(i)| &\leq 2 \max_{i \in \mathcal{I}} |Q(x, \xi^{(i)}) - Q(x', \xi^{(i)})| + \|p_x - p_{x'}\|_1 \max_{i \in \mathcal{I}} |Q(x', \xi^{(i)}) - Q(x^*, \xi^{(i)})| \\ &\leq (2L'_Q + DL'_Q L_p) \|x - x'\|. \end{aligned}$$

Next, we bound $\sup_{i \in \mathcal{I}} |\widehat{Z}_x(i) - \widehat{Z}_{x'}(i)|$. By the representation of the Poisson solution,

$$\begin{aligned} \widehat{Z}_x(i) - \widehat{Z}_{x'}(i) &= \underbrace{\sum_{j \in \mathcal{I}} [(\text{Id} - P_x + \mathbf{1}p_x^\top)^{-1}]_{ij} (e_x(j) - e_{x'}(j))}_E \\ &\quad + \underbrace{\sum_{j \in \mathcal{I}} [(\text{Id} - P_x + \mathbf{1}p_x^\top)^{-1} - (\text{Id} - P_{x'} + \mathbf{1}p_{x'}^\top)^{-1}]_{ij} e_{x'}(j)}_F. \end{aligned}$$

a. We first bound the E term. By the triangle inequality,

$$\begin{aligned} \left| \sum_{j \in \mathcal{I}} [(\text{Id} - P_x + \mathbf{1}p_x^\top)^{-1}]_{ij} (e_x(j) - e_{x'}(j)) \right| &\leq \sum_{j \in \mathcal{I}} \left| [(\text{Id} - P_x + \mathbf{1}p_x^\top)^{-1}]_{ij} \right| |e_x(j) - e_{x'}(j)| \\ &\leq \left(\max_{k \in \mathcal{I}} \sum_{j \in \mathcal{I}} \left| [(\text{Id} - P_x + \mathbf{1}p_x^\top)^{-1}]_{kj} \right| \right) \sup_{\ell \in \mathcal{I}} |e_x(\ell) - e_{x'}(\ell)| \\ &= \|(\text{Id} - P_x + \mathbf{1}p_x^\top)^{-1}\|_{\infty \rightarrow \infty} \sup_{\ell \in \mathcal{I}} |e_x(\ell) - e_{x'}(\ell)| \\ &\leq C_*(2L'_Q + DL'_Q L_p) \|x - x'\|. \end{aligned}$$

Here, the equality uses $\|A\|_{\infty \rightarrow \infty} = \max_k \sum_j |A_{kj}|$, and the last inequality follows from property (B) and the preceding bound on $e_x - e_{x'}$.

b. We next bound the F term. We recall the resolvent identity

$$A^{-1} - B^{-1} = A^{-1}(B - A)B^{-1} \quad \text{with} \quad A := \text{Id} - P_x + \mathbf{1}p_x^\top, \quad B := \text{Id} - P_{x'} + \mathbf{1}p_{x'}^\top.$$

Then,

$$\begin{aligned} &\sup_{i \in \mathcal{I}} \left| \left[\left((\text{Id} - P_x + \mathbf{1}p_x^\top)^{-1} - (\text{Id} - P_{x'} + \mathbf{1}p_{x'}^\top)^{-1} \right) e_{x'} \right]_i \right| \\ &= \left\| \left((\text{Id} - P_x + \mathbf{1}p_x^\top)^{-1} - (\text{Id} - P_{x'} + \mathbf{1}p_{x'}^\top)^{-1} \right) e_{x'} \right\|_\infty \\ &= \left\| (\text{Id} - P_x + \mathbf{1}p_x^\top)^{-1} \left((\text{Id} - P_{x'} + \mathbf{1}p_{x'}^\top) - (\text{Id} - P_x + \mathbf{1}p_x^\top) \right) (\text{Id} - P_{x'} + \mathbf{1}p_{x'}^\top)^{-1} e_{x'} \right\|_\infty \\ &= \left\| (\text{Id} - P_x + \mathbf{1}p_x^\top)^{-1} (P_x - P_{x'} + \mathbf{1}(p_{x'} - p_x)^\top) (\text{Id} - P_{x'} + \mathbf{1}p_{x'}^\top)^{-1} e_{x'} \right\|_\infty. \end{aligned}$$

The chain of equalities continues as

$$\begin{aligned} &= \max_k \left| \sum_j [(\text{Id} - P_x + \mathbf{1}p_x^\top)^{-1}]_{kj} \left[(P_x - P_{x'} + \mathbf{1}(p_{x'} - p_x)^\top) (\text{Id} - P_{x'} + \mathbf{1}p_{x'}^\top)^{-1} e_{x'} \right]_j \right| \\ &\leq \max_k \sum_j \left| [(\text{Id} - P_x + \mathbf{1}p_x^\top)^{-1}]_{kj} \right| \left\| (P_x - P_{x'} + \mathbf{1}(p_{x'} - p_x)^\top) (\text{Id} - P_{x'} + \mathbf{1}p_{x'}^\top)^{-1} e_{x'} \right\|_\infty \\ &= \|(\text{Id} - P_x + \mathbf{1}p_x^\top)^{-1}\|_{\infty \rightarrow \infty} \left\| (P_x - P_{x'} + \mathbf{1}(p_{x'} - p_x)^\top) (\text{Id} - P_{x'} + \mathbf{1}p_{x'}^\top)^{-1} e_{x'} \right\|_\infty \\ &\leq \|(\text{Id} - P_x + \mathbf{1}p_x^\top)^{-1}\|_{\infty \rightarrow \infty} \|P_x - P_{x'} + \mathbf{1}(p_{x'} - p_x)^\top\|_{\infty \rightarrow \infty} \|(\text{Id} - P_{x'} + \mathbf{1}p_{x'}^\top)^{-1}\|_{\infty \rightarrow \infty} \|e_{x'}\|_\infty \\ &\leq \|(\text{Id} - P_x + \mathbf{1}p_x^\top)^{-1}\|_{\infty \rightarrow \infty} (\|P_x - P_{x'}\|_{\infty \rightarrow \infty} + \|p_{x'} - p_x\|_1) \|(\text{Id} - P_{x'} + \mathbf{1}p_{x'}^\top)^{-1}\|_{\infty \rightarrow \infty} \|e_{x'}\|_\infty, \end{aligned}$$

where the last inequality follows from

$$\|P_x - P_{x'} + \mathbf{1}(p_{x'} - p_x)^\top\|_{\infty \rightarrow \infty} \leq \|P_x - P_{x'}\|_{\infty \rightarrow \infty} + \|p_{x'} - p_x\|_1.$$

It follows that the F term is bounded by $2C_*^2 DL'_Q(L_P + L_p)\|x - x'\|$ using properties (B)–(D) and $\|e_{x'}\|_\infty \leq 2DL'_Q$, as proved in (I). Combining the terms E and F gives $\sup_{i \in \mathcal{I}} |\hat{Z}_x(i) - \hat{Z}_{x'}(i)| \leq L_Z\|x - x'\|$.

Finally, we prove that $P_x \hat{Z}_x(i)$ is Lipschitz continuous through

$$\begin{aligned} \left| (P_x \hat{Z}_x)(i) - (P_{x'} \hat{Z}_{x'})(i) \right| &\leq \sum_j |P_x(i, j) - P_{x'}(i, j)| |\hat{Z}_x(j)| + \sum_j P_{x'}(i, j) |\hat{Z}_x(j) - \hat{Z}_{x'}(j)| \\ &\leq \left(\sum_j |P_x(i, j) - P_{x'}(i, j)| \right) \sup_j |\hat{Z}_x(j)| + \sup_j |\hat{Z}_x(j) - \hat{Z}_{x'}(j)| \\ &\leq (L_P B_Z + L_Z) \|x - x'\|. \quad \square \end{aligned}$$

We are now ready to present a proof for Proposition 1.

Proof of Proposition 1: We prove the following equivalent claim:

$$\mathbb{E} \left[-2 \sum_{k=0}^{n-1} \alpha_{k+1} e_{k+1} \right] \leq 2(L'_Q + L_f)(L_P B_Z + L_Z) \sum_{k=0}^{n-1} \alpha_{k+1}^2 + 4B_Z \alpha_1.$$

Let $\mathcal{F}_k := \sigma(\{(x_t, I_t) : 0 \leq t \leq k\})$. Recall the Poisson equation in Lemma EC.7 has $e_{k+1} = \hat{Z}_{x_k}(I_{k+1}) - (P_{x_k} \hat{Z}_{x_k})(I_{k+1})$. Since $I_{k+1} \sim P_{x_k}(I_k, \cdot)$ conditionally on $\mathcal{F}_k$, $\mathbb{E}[\hat{Z}_{x_k}(I_{k+1}) | \mathcal{F}_k] = (P_{x_k} \hat{Z}_{x_k})(I_k)$. The tower property gives

$$\mathbb{E} \left[-2 \sum_{k=0}^{n-1} \alpha_{k+1} e_{k+1} \right] = -2 \sum_{k=0}^{n-1} \mathbb{E} \left[\alpha_{k+1} ((P_{x_k} \hat{Z}_{x_k})(I_k) - (P_{x_k} \hat{Z}_{x_k})(I_{k+1})) \right]. \quad (\text{I})$$

Add and subtract $(P_{x_{k+1}} \hat{Z}_{x_{k+1}})(I_{k+1})$ inside (I) and split the sum into

$$\underbrace{-2 \sum_{k=0}^{n-1} \alpha_{k+1} ((P_{x_k} \hat{Z}_{x_k})(I_k) - (P_{x_{k+1}} \hat{Z}_{x_{k+1}})(I_{k+1}))}_A - 2 \sum_{k=0}^{n-1} \alpha_{k+1} ((P_{x_{k+1}} \hat{Z}_{x_{k+1}})(I_{k+1}) - (P_{x_k} \hat{Z}_{x_k})(I_{k+1})) \quad (\text{II})$$

For the A term in (II), we shift the index to obtain that $A = -2\alpha_1(P_{x_0} \hat{Z}_{x_0})(I_0) + 2\alpha_n(P_{x_n} \hat{Z}_{x_n})(I_n) - 2\sum_{k=1}^{n-1} (\alpha_{k+1} - \alpha_k)(P_{x_k} \hat{Z}_{x_k})(I_k)$. Using $|(P_x \hat{Z}_x)(i)| \leq B_Z$ from Lemma EC.7 and the nonincreasing step sizes, we have

$$\mathbb{E}[A] \leq 2(\alpha_1 + \alpha_n)B_Z + 2 \sum_{k=1}^{n-1} (\alpha_k - \alpha_{k+1})B_Z = 4B_Z \alpha_1. \quad (\text{III})$$

For the B sum in (II), Lemma EC.7 gives $|(P_{x_{k+1}} \hat{Z}_{x_{k+1}})(I_{k+1}) - (P_{x_k} \hat{Z}_{x_k})(I_{k+1})| \leq (L_P B_Z + L_Z)\|x_{k+1} - x_k\|$. Now recall that $x_{k+1} = \Pi_X(x_k - \alpha_{k+1} Z_{x_k}(I_{k+1}))$. Since $x_k \in X$, we have

$\Pi_X(x_k) = x_k$. Then, by the nonexpansiveness of the projection operator, we have $\|x_{k+1} - x_k\|_1 = \|\Pi_X(x_k - \alpha_k Z_{x_k}(I_{k+1})) - \Pi_X(x_k)\|_1 \leq \alpha_k \|Z_{x_k}(I_{k+1})\|_1 \leq \alpha_k (L_f + L'_Q)$, where we use the bound $\|Z_x(i)\| \leq L_f + L'_Q$ established in Lemma 3. It follows that

$$\mathbb{E}[B] \leq 2(L_P B_Z + L_Z) \sum_{k=0}^{n-1} \alpha_{k+1} \mathbb{E}\|x_{k+1} - x_k\| \leq 2(L_f + L'_Q)(L_P B_Z + L_Z) \sum_{k=0}^{n-1} \alpha_{k+1}^2 \quad (\text{IV})$$

Finally, combining (III) and (IV) completes the proof. $\square$

EC.4.5. Proof of Theorem 1

Proof: By nonexpansiveness of the Euclidean projection and that $x^* \in \mathcal{X}$, we have

$$\|x_k - x^*\|^2 \leq \|x_{k-1} - \alpha_k Z_{x_{k-1}}(I_k) - x^*\|^2 = \|x_{k-1} - x^*\|^2 - 2\alpha_k \cdot Z_{x_{k-1}}(I_k)^\top (x_{k-1} - x^*) + \alpha_k^2 \|Z_{x_{k-1}}(I_k)\|^2.$$

Recall, from the derivation before Proposition 1, that $Z_{x_{k-1}}(I_k)^\top (x_{k-1} - x^*) \geq F_w(x_{k-1}) - F_w(x^*) + e_k$. Together with $\|Z_{x_{k-1}}(I_k)\| \leq L_f + L'_Q$ from Lemma 3, this gives $2\alpha_k (F_w(x_{k-1}) - F_w^*) \leq \|x_{k-1} - x^*\|^2 - \|x_k - x^*\|^2 - 2\alpha_k e_k + \alpha_k^2 (L_f + L'_Q)^2$. To pass to the updated iterate, note that $\|x_k - x_{k-1}\| \leq \alpha_k (L_f + L'_Q)$ and F_w is $(L_f + L'_Q)$ -Lipschitz on $\mathcal{X}$ by Lemma 3. Then, $F_w(x_k) - F_w(x_{k-1}) \leq (L_f + L'_Q)\|x_k - x_{k-1}\| \leq \alpha_k (L_f + L'_Q)^2$. Combining these inequalities yields

$$2\alpha_k (F_w(x_k) - F_w^*) \leq \|x_{k-1} - x^*\|^2 - \|x_k - x^*\|^2 - 2\alpha_k e_k + 3\alpha_k^2 (L_f + L'_Q)^2.$$

Summing this inequality over $k = 1, \dots, n$ gives

$$\begin{aligned} 2 \sum_{k=1}^n \alpha_k (F_w(x_k) - F_w^*) &\leq \sum_{k=1}^n (\|x_{k-1} - x^*\|^2 - \|x_k - x^*\|^2) - 2 \sum_{k=1}^n \alpha_k e_k + 3(L_f + L'_Q)^2 \sum_{k=1}^n \alpha_k^2 \\ &= \|x_0 - x^*\|^2 - \|x_n - x^*\|^2 - 2 \sum_{k=1}^n \alpha_k e_k + 3(L_f + L'_Q)^2 \sum_{k=1}^n \alpha_k^2. \end{aligned}$$

Taking expectations and dropping the nonpositive term $-\mathbb{E}[\|x_n - x^*\|^2]$, we obtain

$$\begin{aligned} 2 \sum_{k=1}^n \alpha_k \mathbb{E}[F_w(x_k) - F_w^*] &\leq \mathbb{E}[\|x_0 - x^*\|^2] + \mathbb{E}\left[-2 \sum_{k=1}^n \alpha_k e_k\right] + 3(L_f + L'_Q)^2 \sum_{k=1}^n \alpha_k^2 \\ &\leq \mathbb{E}[\|x_0 - x^*\|^2] + C_0 + (C_1 + 3(L_f + L'_Q)^2) \sum_{k=1}^n \alpha_k^2, \end{aligned}$$

where the last inequality follows from Proposition 1. But the convexity of F_w and the definition of $\bar{x}_n$ imply that $\mathbb{E}[F_w(\bar{x}_n) - F_w(x^*)] \leq \frac{1}{\sum_{k=1}^n \alpha_k} \sum_{k=1}^n \alpha_k \mathbb{E}[F_w(x_k) - F_w(x^*)]$. Applying the upper bound for $\sum_{k=1}^n \alpha_k \mathbb{E}[F_w(x_k) - F_w^*]$ we just derived above produces

$$\mathbb{E}[F_w(\bar{x}_n) - F_w(x^*)] \leq \frac{\mathbb{E}[\|x_0 - x^*\|^2] + C_0 + (C_1 + 3(L_f + L'_Q)^2) \sum_{k=1}^n \alpha_k^2}{2 \sum_{k=1}^n \alpha_k}.$$

For the step-size average $\bar{x}_n^{\text{ss}} := (\sum_{k=1}^n \alpha_k x_k) / (\sum_{k=1}^n \alpha_k)$ with $\alpha_k = O(1/\sqrt{k})$, we have $\sum_{k=1}^n \alpha_k = \Theta(\sqrt{n})$ and $\sum_{k=1}^n \alpha_k^2 = O(1 + \log n)$. Since $L_P = O(1/w)$, $B_Z = O(C_*)$, $L_Z = O(C_* + \frac{C_*^2}{w})$, $C_1 = O(C_* + \frac{C_*^2}{w})$, and $C_0 = O(C_*)$, plugging these bounds into the above inequality leads to the claimed convergence rate.

Furthermore, for the linear averages $\bar{x}_n^{\text{lin}} := (\sum_{k=1}^n kx_k) / (\sum_{k=1}^n k)$, multiplying the preceding one-step inequality by $k/(2\alpha_k)$ gives

$$k(F_w(x_k) - F_w^*) \leq \frac{k}{2\alpha_k} (\|x_{k-1} - x^*\|^2 - \|x_k - x^*\|^2) - ke_k + \frac{3}{2}k\alpha_k(L_f + L'_Q)^2.$$

We bound the three terms after summing over k . First, since $\alpha_k = a/\sqrt{k}$, the sequence k/α_k is increasing. Summation by parts yields

$$\begin{aligned} \sum_{k=1}^n \frac{k}{2\alpha_k} (\|x_{k-1} - x^*\|^2 - \|x_k - x^*\|^2) &= \frac{\|x_0 - x^*\|^2}{2\alpha_1} - \frac{n\|x_n - x^*\|^2}{2\alpha_n} + \sum_{k=1}^{n-1} \left(\frac{k+1}{2\alpha_{k+1}} - \frac{k}{2\alpha_k} \right) \|x_k - x^*\|^2 \\ &\leq \frac{D^2}{2\alpha_1} + D^2 \sum_{k=1}^{n-1} \left(\frac{k+1}{2\alpha_{k+1}} - \frac{k}{2\alpha_k} \right) = \frac{D^2 n}{2\alpha_n}, \end{aligned}$$

where we used $\|x_k - x^*\| \leq D$ and dropped the nonpositive terminal term. Second, the Poisson equation and the tower property give $\mathbb{E} \sum_{k=1}^n ke_k = \mathbb{E} \sum_{k=1}^n k((P_{x_{k-1}} \hat{Z}_{x_{k-1}})(I_{k-1}) - (P_{x_{k-1}} \hat{Z}_{x_{k-1}})(I_k))$. Shifting the index and collecting terms, we obtain

$$\begin{aligned} \mathbb{E} \sum_{k=1}^n ke_k &= \mathbb{E} \left[(P_{x_0} \hat{Z}_{x_0})(I_0) - n(P_{x_{n-1}} \hat{Z}_{x_{n-1}})(I_n) + \sum_{k=1}^{n-1} (P_{x_{k-1}} \hat{Z}_{x_{k-1}})(I_k) \right. \\ &\quad \left. + \sum_{k=1}^{n-1} (k+1)((P_{x_k} \hat{Z}_{x_k})(I_k) - (P_{x_{k-1}} \hat{Z}_{x_{k-1}})(I_k)) \right]. \end{aligned}$$

Then, Lemma EC.7 and $\|x_k - x_{k-1}\| \leq \alpha_k(L_f + L'_Q)$ imply $|\mathbb{E} \sum_{k=1}^n ke_k| \leq 2nB_Z + (L_f + L'_Q)(L_P B_Z + L_Z) \sum_{k=1}^{n-1} (k+1)\alpha_k$. Combining these bounds and applying convexity, we conclude that

$$\begin{aligned} \mathbb{E}[F_w(\bar{x}_n^{\text{lin}})] - F_w^* &\leq \frac{2}{n(n+1)} \sum_{k=1}^n k \mathbb{E}[F_w(x_k) - F_w^*] \\ &\leq \frac{2}{n(n+1)} \left[\frac{D^2 n}{2\alpha_n} + 2nB_Z + \frac{3}{2}(L_f + L'_Q)^2 \sum_{k=1}^n k\alpha_k \right. \\ &\quad \left. + (L_f + L'_Q)(L_P B_Z + L_Z) \sum_{k=1}^{n-1} (k+1)\alpha_k \right]. \end{aligned}$$

For $\alpha_k = a/\sqrt{k}$, $n/\alpha_n = O(n^{3/2})$, $\sum_{k=1}^n k\alpha_k = O(n^{3/2})$, and $\sum_{k=1}^{n-1} (k+1)\alpha_k = O(n^{3/2})$. Using $B_Z = O(C_*)$ and $L_P B_Z + L_Z = O(C_* + C_*^2/w)$, we arrive at

$$\mathbb{E}[F_w(\bar{x}_n^{\text{lin}})] - F_w^* = O\left(\left(C_* + \frac{C_*^2}{w}\right)n^{-1/2}\right). \quad \square$$

EC.4.6. Proof of Proposition 2

Proof of Proposition 2: Fix $x^* \in \arg \min_{x \in \mathcal{X}} F(x)$ and define $e_k := e_{x_{k-1}, w_{k-1}}(I_k)$ with

$$e_{x,w}(i) := Q(x, \xi^{(i)}) - Q(x^*, \xi^{(i)}) - \sum_j p_{x,w}(j) (Q(x, \xi^{(j)}) - Q(x^*, \xi^{(j)}))$$

for all $i \in \mathcal{I}$. Let $\widehat{Z}_{x,w}$ denote the corresponding centered solution to the Poisson equation such that $\widehat{Z}_{x,w} - P_{x,w} \widehat{Z}_{x,w} = e_{x,w}$ and $p_{x,w}^\top \widehat{Z}_{x,w} = 0$, and $p_{x,w}$ denote the stationary distribution with respect to incumbent x and temperature parameter w with $p_{x,w}(i) := \exp(Q(x, \xi^{(i)})/w) / \sum_{j \in \mathcal{I}} \exp(Q(x, \xi^{(j)})/w)$. We first derive a useful lemma. Note that the Lipschitz continuity of $p_{x,w}$ in (x, w) holds for a general state space $\mathcal{I}$.

LEMMA EC.8 (Lipschitz continuity in the iterate and temperature). *For $Q(\cdot, \cdot)$ satisfying the Lipschitz continuity in Lemma 1, suppose that the proposal kernel is fixed and the acceptance distribution is either MH or Barker. Then, for all $x, y \in \mathcal{X}$ and $w, v > 0$, it holds that*

$$\|p_{x,w} - p_{y,v}\|_1 \leq \frac{2L'_Q}{\max\{w, v\}} \|x - y\| + 2L'_\xi \text{diam}(\Xi) \left| \frac{1}{w} - \frac{1}{v} \right|.$$

Moreover, the MH transition kernel satisfies

$$\|P_{x,w} - P_{y,v}\|_{\infty \rightarrow \infty} \leq \frac{4L'_Q}{\max\{w, v\}} \|x - y\| + 2L'_\xi \text{diam}(\Xi) \left| \frac{1}{w} - \frac{1}{v} \right|.$$

For the Barker transition kernel, the bound improves to

$$\|P_{x,w} - P_{y,v}\|_{\infty \rightarrow \infty} \leq \frac{L'_Q}{\max\{w, v\}} \|x - y\| + \frac{L'_\xi \text{diam}(\Xi)}{2} \left| \frac{1}{w} - \frac{1}{v} \right|.$$

Proof of Lemma EC.8: We already proved in Lemma EC.6 that $\|p_{x,w} - p_{y,w}\|_1 \leq (2L'_Q/w) \|x - y\|$. Fix a state $i_0 \in \mathcal{I}$. Since the recourse costs have bounded variation across scenarios, differentiation under the integral is justified. Differentiating with respect to the inverse temperature w^{-1} gives

$$\begin{aligned} \int_{\mathcal{I}} \left| \frac{\partial p_{y,w}(i)}{\partial(1/w)} \right| \nu(di) &= \int_{\mathcal{I}} p_{y,w}(i) \left| Q(y, \xi^{(i)}) - \int_{\mathcal{I}} p_{y,w}(j) Q(y, \xi^{(j)}) \nu(dj) \right| \nu(di) \\ &= \int_{\mathcal{I}} p_{y,w}(i) \left| Q(y, \xi^{(i)}) - Q(y, \xi^{(i_0)}) - \int_{\mathcal{I}} p_{y,w}(j) (Q(y, \xi^{(j)}) - Q(y, \xi^{(i_0)})) \nu(dj) \right| \nu(di) \\ &\leq \int_{\mathcal{I}} p_{y,w}(i) |Q(y, \xi^{(i)}) - Q(y, \xi^{(i_0)})| \nu(di) \\ &\quad + \int_{\mathcal{I}} p_{y,w}(i) \int_{\mathcal{I}} p_{y,w}(j) |Q(y, \xi^{(j)}) - Q(y, \xi^{(i_0)})| \nu(dj) \nu(di) \\ &= 2 \int_{\mathcal{I}} p_{y,w}(i) |Q(y, \xi^{(i)}) - Q(y, \xi^{(i_0)})| \nu(di) \\ &\leq 2 \sup_{i \in \mathcal{I}} |Q(y, \xi^{(i)}) - Q(y, \xi^{(i_0)})| \int_{\mathcal{I}} p_{y,w}(i) \nu(di) = 2 \sup_{i \in \mathcal{I}} |Q(y, \xi^{(i)}) - Q(y, \xi^{(i_0)})|. \end{aligned}$$

Then, integrating between $1/w$ and $1/v$ yields

$$\|p_{y,w} - p_{y,v}\|_1 \leq 2 \sup_{i \in \mathcal{I}} |Q(y, \xi^{(i)}) - Q(y, \xi^{(i_0)})| \cdot \left| \frac{1}{w} - \frac{1}{v} \right| \leq 2L'_\xi \text{diam}(\Xi) \left| \frac{1}{w} - \frac{1}{v} \right|.$$

In addition, $\|p_{x,w} - p_{y,v}\|_1 \leq \|p_{x,w} - p_{y,w}\|_1 + \|p_{y,w} - p_{y,v}\|_1 \leq (2L'_Q/w)\|x - y\| + 2L'_\xi \text{diam}(\Xi) |1/w - 1/v|$ and $\|p_{x,w} - p_{y,v}\|_1 \leq \|p_{x,w} - p_{x,v}\|_1 + \|p_{x,v} - p_{y,v}\|_1 \leq (2L'_Q/v)\|x - y\| + 2L'_\xi \text{diam}(\Xi) |1/w - 1/v|$. Thus, $\|p_{x,w} - p_{y,v}\|_1 \leq (2L'_Q/\max\{w, v\})\|x - y\| + 2L'_\xi \text{diam}(\Xi) |1/w - 1/v|$.

We next bound the transition kernels and the proof is similar to those of Lemmas [EC.3](#) and [EC.4](#). Write the logarithmic acceptance ratio as $\alpha_{x,w}(i, j) := \frac{Q(x, \xi^{(j)}) - Q(x, \xi^{(i)})}{w} + \log q(i | j) - \log q(j | i)$. Lipschitz continuity of Q therefore gives $|\alpha_{x,w}(i, j) - \alpha_{y,v}(i, j)| \leq (2L'_Q/\max\{w, v\})\|x - y\| + L'_\xi \text{diam}(\Xi) |1/w - 1/v|$. The remaining proof is the same, so we will omit it here. $\square$

Proof of Proposition 2 (continued): Furthermore, we adapt the same resolvent argument to $\sup_i |(P_{x,w} \hat{Z}_{x,w})(i) - (P_{y,v} \hat{Z}_{y,v})(i)|$. The definition of $e_{x,w}$ and $|Q(y, \xi^{(i)}) - Q(x^*, \xi^{(i)})| \leq DL'_Q$ imply (i) $\|e_{x,w}\|_\infty \leq 2DL'_Q$ and (ii) $\|e_{x,w} - e_{y,v}\|_\infty \leq 2L'_Q\|x - y\| + DL'_Q\|p_{x,w} - p_{y,v}\|_1$, thus we have for

$$\begin{aligned} \hat{Z}_{x,w} - \hat{Z}_{y,v} &= (\text{Id} - P_{x,w} + \mathbf{1}p_{x,w}^\top)^{-1}(e_{x,w} - e_{y,v}) \\ &\quad + \left((\text{Id} - P_{x,w} + \mathbf{1}p_{x,w}^\top)^{-1} - (\text{Id} - P_{y,v} + \mathbf{1}p_{y,v}^\top)^{-1} \right) e_{y,v} \\ &= (\text{Id} - P_{x,w} + \mathbf{1}p_{x,w}^\top)^{-1}(e_{x,w} - e_{y,v}) \\ &\quad + (\text{Id} - P_{x,w} + \mathbf{1}p_{x,w}^\top)^{-1}(P_{x,w} - P_{y,v} + \mathbf{1}(p_{y,v} - p_{x,w})^\top) \cdot (\text{Id} - P_{y,v} + \mathbf{1}p_{y,v}^\top)^{-1} e_{y,v}, \end{aligned}$$

where the second equality uses the resolvent identity $A^{-1} - B^{-1} = A^{-1}(B - A)B^{-1}$ with $A := \text{Id} - P_{x,w} + \mathbf{1}p_{x,w}^\top$ and $B := \text{Id} - P_{y,v} + \mathbf{1}p_{y,v}^\top$. It follows from Assumption 4 that $\|\hat{Z}_{x,w} - \hat{Z}_{y,v}\|_\infty \leq C_* \|e_{x,w} - e_{y,v}\|_\infty + 2C_*^2 DL'_Q (\|P_{x,w} - P_{y,v}\|_{\infty \rightarrow \infty} + \|p_{x,w} - p_{y,v}\|_1)$. Finally, the Poisson equation gives $P_{x,w} \hat{Z}_{x,w} = \hat{Z}_{x,w} - e_{x,w}$. Hence,

$$\begin{aligned} &\sup_i |(P_{x,w} \hat{Z}_{x,w})(i) - (P_{y,v} \hat{Z}_{y,v})(i)| \\ &\leq \|e_{x,w} - e_{y,v}\|_\infty + \|\hat{Z}_{x,w} - \hat{Z}_{y,v}\|_\infty \\ &\leq (C_* + 1) \|e_{x,w} - e_{y,v}\|_\infty + 2C_*^2 DL'_Q (\|P_{x,w} - P_{y,v}\|_{\infty \rightarrow \infty} + \|p_{x,w} - p_{y,v}\|_1) \\ &\leq 2(C_* + 1) L'_Q \|x - y\| + (C_* + 1 + 2C_*^2) DL'_Q \|p_{x,w} - p_{y,v}\|_1 + 2C_*^2 DL'_Q \|P_{x,w} - P_{y,v}\|_{\infty \rightarrow \infty} \\ &\leq C_*^2 \left[(4L'_Q + 16D(L'_Q)^2) \frac{\|x - y\|}{\min\{w, v\}} + 12DL'_Q L'_\xi \text{diam}(\Xi) \left| \frac{1}{w} - \frac{1}{v} \right| \right], \end{aligned}$$

where we use $C_* \geq 1$ and $\min\{w, v\} \leq 1$ for sufficiently large n considered. We now adapt the proof of Theorem 1 with $e_k := e_{x_{k-1}, w_{k-1}}(I_k)$. We obtain similarly that

$$\begin{aligned} \mathbb{E} \sum_{k=1}^n k e_k &= \mathbb{E} \left[(P_{x_0, w_0} \hat{Z}_{x_0, w_0})(I_0) - n(P_{x_{n-1}, w_{n-1}} \hat{Z}_{x_{n-1}, w_{n-1}})(I_n) + \sum_{k=1}^{n-1} (P_{x_{k-1}, w_{k-1}} \hat{Z}_{x_{k-1}, w_{k-1}})(I_k) \right. \\ &\quad \left. + \sum_{k=1}^{n-1} (k+1) ((P_{x_k, w_k} \hat{Z}_{x_k, w_k})(I_k) - (P_{x_{k-1}, w_{k-1}} \hat{Z}_{x_{k-1}, w_{k-1}})(I_k)) \right]. \end{aligned}$$

Collect the results to use. By Lemma EC.7, $\sup_i |(P_{x,w}\widehat{Z}_{x,w})(i)| = \mathcal{O}(C_*)$ and we just proved a bound for $\sup_i |(P_{x,w}\widehat{Z}_{x,w})(i) - (P_{y,v}\widehat{Z}_{y,v})(i)|$:

$$\sup_i |(P_{x_k,w_k}\widehat{Z}_{x_k,w_k})(i) - (P_{x_{k-1},w_{k-1}}\widehat{Z}_{x_{k-1},w_{k-1}})(i)| = \mathcal{O}\left(C_*^2 \left[\frac{\|x_k - x_{k-1}\|}{w_k} + \left| \frac{1}{w_k} - \frac{1}{w_{k-1}} \right| \right] \right),$$

where we also use $w_k \leq w_{k-1}$. It follows from $\|x_k - x_{k-1}\| \leq \alpha_k(L_f + L'_Q)$ that

$$\left| \mathbb{E} \sum_{k=1}^n k e_k \right| = 2n\mathcal{O}(C_*) + \mathcal{O}\left(C_*^2 \sum_{k=1}^{n-1} (k+1) \left[\frac{\alpha_k}{w_k} + \left| \frac{1}{w_k} - \frac{1}{w_{k-1}} \right| \right] \right)$$

We next account for the approximation error. By the definition of e_k and the subgradient inequality, we have

$$\begin{aligned} Z_{x_{k-1}}(I_k)^\top (x_{k-1} - x^*) &\geq f(x_{k-1}) - f(x^*) + \sum_j p_{x_{k-1},w_{k-1}}(j) (Q(x_{k-1}, \xi^{(j)}) - Q(x^*, \xi^{(j)})) + e_k \\ &\geq F(x_{k-1}) - F^* - \left[\max_i Q(x_{k-1}, \xi^{(i)}) - \sum_j p_{x_{k-1},w_{k-1}}(j) Q(x_{k-1}, \xi^{(j)}) \right] + e_k. \end{aligned}$$

Furthermore, because $\Delta_* \leq \widetilde{\Delta}_k$ for all sufficiently large k , there exists a $k_0 \in \mathbb{N}_+$ such that $p_{x_{k-1},w_{k-1}}(\mathcal{I} \setminus I^*(x_{k-1})) \leq \frac{|\mathcal{I}| - |I^*(x_{k-1})|}{|I^*(x_{k-1})|} e^{-\Delta_*/w_{k-1}} \leq |\mathcal{I}| e^{-\Delta_*/w_{k-1}}$ for all $k \geq k_0$. Recall that $\Delta := \sup_x \max_{i,j} |Q(x, \xi^{(i)}) - Q(x, \xi^{(j)})| < \infty$, so for finitely many $k < k_0$, we use the crude bound $\max_i Q(x_{k-1}, \xi^{(i)}) - \sum_j p_{x_{k-1},w_{k-1}}(j) Q(x_{k-1}, \xi^{(j)}) \leq \Delta$, while for all $k \geq k_0$, we have $Z_{x_{k-1}}(I_k)^\top (x_{k-1} - x^*) \geq F(x_{k-1}) - F^* - \Delta |\mathcal{I}| e^{-\Delta_*/w_{k-1}} + e_k$. Following the projection argument as in the proof of Theorem 1, for $k \geq k_0$, we obtain

$$2\alpha_k(F(x_k) - F^*) \leq \|x_{k-1} - x^*\|^2 - \|x_k - x^*\|^2 + 3\alpha_k^2(L_f + L'_Q)^2 - 2\alpha_k e_k + 2\alpha_k \Delta |\mathcal{I}| e^{-\Delta_*/w_{k-1}}.$$

For the finitely many $k < k_0$, the same argument with the crude bound above contributes only a constant $O(1)$ independent of n . We multiply the above inequality by $k/(2\alpha_k)$ on both sides and apply convexity. In the proof of Theorem 1, we have shown that $\sum_{k=1}^n (k/(2\alpha_k)) (\|x_{k-1} - x^*\|^2 - \|x_k - x^*\|^2) \leq (D^2 n)/(2\alpha_n)$. Therefore,

$$\begin{aligned} \mathbb{E}[F(\bar{x}_n)] - F^* &\leq \frac{2}{n(n+1)} \sum_{k=1}^n k \mathbb{E}[F(x_k) - F^*] \\ &\leq \frac{2}{n(n+1)} \left[\frac{D^2 n}{2\alpha_n} + \frac{3}{2} (L_f + L'_Q)^2 \sum_{k=1}^n k \alpha_k + \left| \mathbb{E} \sum_{k=1}^n k e_k \right| + \Delta |\mathcal{I}| \sum_{k=k_0}^n k e^{-\Delta_*/w_{k-1}} + O(1) \right]. \end{aligned}$$

Now let $\alpha_k = a/\sqrt{k}$ and $w_k = (2\Delta_*)/(\log(k+2) + 2\log|\mathcal{I}|)$. Then $\alpha_k/w_k = O(k^{-1/2}(\log k + \log|\mathcal{I}|))$ and $|1/w_k - 1/w_{k-1}| = O(k^{-1})$. Hence,

$$\left| \mathbb{E} \sum_{k=1}^n k e_k \right| = O\left(C_* n + C_*^2 n^{3/2}(\log n + \log|\mathcal{I}|)\right).$$

Moreover, $|\mathcal{I}|e^{-\Delta^*/w_{k-1}} = (k+1)^{-1/2}$, so $\Delta|\mathcal{I}|\sum_{k=k_0}^n ke^{-\Delta^*/w_{k-1}} = O(n^{3/2})$. Since $n/\alpha_n = O(n^{3/2})$ and $\sum_{k=1}^n k\alpha_k = O(n^{3/2})$, it follows that

$$0 \leq \mathbb{E}[F(\bar{x}_n)] - F^* \leq O\left(\frac{1 + C_*^2(\log n + \log |\mathcal{I}|)}{\sqrt{n}} + \frac{C_*}{n}\right) = O\left(\frac{1 + C_*^2 \log(n|\mathcal{I}|)}{\sqrt{n}}\right). \quad \square$$

EC.4.7. General Mellowmax Approximation

PROPOSITION EC.1. *Denote by $(\mathcal{I}, d_{\mathcal{I}})$ a compact metric space with diameter $D_{\mathcal{I}} := \sup\{d_{\mathcal{I}}(I, J) : I, J \in \mathcal{I}\}$ for the states. Suppose that there exists a constant $L_{\mathcal{I}} < \infty$, such that $|\widehat{Q}_x(I) - \widehat{Q}_x(J)| \leq L_{\mathcal{I}} d_{\mathcal{I}}(I, J)$ for all $x \in \mathcal{X}$ and $I, J \in \mathcal{I}$. Suppose further that there exist constants $c_{\nu} > 0$ and $d_I > 0$ such that $\nu(B_{\mathcal{I}}(I, r)) \geq c_{\nu} (r/D_{\mathcal{I}})^{d_I}$ for all $I \in \mathcal{I}$ and $0 < r \leq D_{\mathcal{I}}$. Then, for every $x \in \mathcal{X}$, it holds that*

$$0 \leq \sup_{I \in \mathcal{I}} \widehat{Q}(x, \xi_I) - \mathcal{M}_w^{\nu}(x) \leq \inf_{0 < r \leq D_{\mathcal{I}}} \left\{ L_I r + w \log \frac{1}{c_{\nu}} + d_I w \log \frac{D_{\mathcal{I}}}{r} \right\}.$$

In particular, if $d_I w / L_{\mathcal{I}} \leq D_{\mathcal{I}}$ (which always holds for a sufficiently small w), then

$$\sup_{I \in \mathcal{I}} \widehat{Q}(x, \xi_I) - \mathcal{M}_w^{\nu}(x) \leq w \left(\log \frac{1}{c_{\nu}} + d_I \left[1 + \log \left(\frac{L_{\mathcal{I}} D_{\mathcal{I}}}{d_I w} \right) \right] \right).$$

Proof: We fix $x \in \mathcal{X}$ and write $Q^*(x) = \sup_{I \in \mathcal{I}} \widehat{Q}(x, \xi_I)$. Since $\mathcal{I}$ is compact and $I \mapsto \widehat{Q}(x, \xi_I)$ is Lipschitz continuous, there exists an $I_x^* \in \mathcal{I}$ that attains the supremum. Now because ν is a probability measure, we have $\mathcal{M}_w^{\nu}(x) = w \log \int_{\mathcal{I}} \exp(\widehat{Q}(x, \xi_I)/w) \nu(dI) \leq w \log \int_{\mathcal{I}} \exp(Q^*(x)/w) \nu(dI) = Q^*(x)$. This proves that $Q^*(x) - \mathcal{M}_w^{\nu}(x) \geq 0$.

Next, we fix any $r \in (0, D_{\mathcal{I}}]$. For every $I \in B_{\mathcal{I}}(I_x^*, r)$, the Lipschitz property gives $\widehat{Q}(x, \xi_I) \geq \widehat{Q}(x, \xi_{I_x^*}) - L_{\mathcal{I}} d_{\mathcal{I}}(I, I_x^*) \geq Q^*(x) - L_{\mathcal{I}} r$. It follows that

$$\int_{\mathcal{I}} \exp\left(\frac{\widehat{Q}(x, \xi_I)}{w}\right) \nu(dI) \geq \int_{B_{\mathcal{I}}(I_x^*, r)} \exp\left(\frac{\widehat{Q}(x, \xi_I)}{w}\right) \nu(dI) \geq \exp\left(\frac{Q^*(x) - L_{\mathcal{I}} r}{w}\right) \nu(B_{\mathcal{I}}(I_x^*, r)).$$

Taking logarithms and multiplying by w on both sides yield $\mathcal{M}_w^{\nu}(x) \geq Q^*(x) - L_{\mathcal{I}} r + w \log \nu(B_{\mathcal{I}}(I_x^*, r))$. But $w \log \nu(B_{\mathcal{I}}(I_x^*, r)) \geq w \log [c_{\nu} (r/D_{\mathcal{I}})^{d_I}] = -w \log(1/c_{\nu}) - d_I w \log(D_{\mathcal{I}}/r)$. It follows that $Q^*(x) - \mathcal{M}_w^{\nu}(x) \leq L_{\mathcal{I}} r + w \log(1/c_{\nu}) + d_I w \log(D_{\mathcal{I}}/r)$. Since this inequality holds for every $r \in (0, D_{\mathcal{I}}]$, taking the infimum over r proves the first claim.

For $L_{\mathcal{I}} > 0$, we find a smallest upper bound by choosing r . Under the condition $d_I w / L_{\mathcal{I}} \leq D_{\mathcal{I}}$, the unconstrained minimizer $r^* = d_I w / L_{\mathcal{I}}$ is feasible. Substitution gives the desired result. $\square$

EC.4.8. Lower Bound of $\|dp_x^w/d\nu\|_{\infty}^{-1}$ in Example 9

Recall that, because $\mathcal{I} = \mathbb{S}^{d_{\xi}-1}$, we take B as $B_{\mathbb{S}}(I_x^*, \theta) := \{x \in \mathbb{S}^{d_{\xi}-1} : \angle(x, e_1) \leq \theta\}$, where $I_x^* \in \arg \max_I \widehat{Q}(I)$ and we assume WLOG that $I_x^* = e_1$ because $\mathcal{I}$ is rotationally invariant. Because $I \mapsto \widehat{Q}(I)$ is $L'_{\xi} \|A\|_{2 \rightarrow 2}$ -Lipschitz, for every $J \in B_{\mathbb{S}}(I_x^*, \theta)$, it holds that $\widehat{Q}(J) \geq \widehat{Q}(I_x^*) -$

$(2 \sin \frac{\theta}{2}) L'_\xi \|A\|_{2 \rightarrow 2} \geq \widehat{Q}(I_x^*) - L'_\xi \|A\|_{2 \rightarrow 2} \theta$. Then, $\|dp_x^w/d\nu\|_\infty^{-1} = \int_{\mathcal{I}} \exp\left\{(\widehat{Q}(J) - \widehat{Q}(I_x^*))/w\right\} \nu(dJ) \geq \nu(B_{\mathbb{S}}(I_x^*, \theta)) \exp\left\{-L'_\xi \|A\|_{2 \rightarrow 2} \theta/w\right\}$. We claim that the uniform measure of a spherical cap satisfies

$$\nu(B_{\mathbb{S}}(I_x^*, \theta)) \geq c_{d_\xi} \theta^{d_\xi-1} \quad \forall 0 < \theta \leq \frac{\pi}{2},$$

where $c_{d_\xi} > 0$ depends only on d_ξ .

Proof of the claim: Every $x \in \mathbb{S}^{d_\xi-1}$ can be written as $x = (\cos t, \sin t \omega)$, $t \in [0, \pi]$, $\omega \in \mathbb{S}^{d_\xi-2}$, where $t = \angle(x, e_1)$. In these spherical coordinates, the surface element is $d\sigma_{d_\xi-1}(x) = \sin^{d_\xi-2}(t) dt d\sigma_{d_\xi-2}(\omega)$. Because ν is uniform on $\mathcal{I}$,

$$\nu(B_S(I_x^*, \theta)) = \frac{\int_{\mathbb{S}^{d_\xi-2}} \int_0^\theta \sin^{d_\xi-2}(t) dt d\sigma_{d_\xi-2}(\omega)}{\int_{\mathbb{S}^{d_\xi-2}} \int_0^\pi \sin^{d_\xi-2}(t) dt d\sigma_{d_\xi-2}(\omega)} = \frac{\int_0^\theta \sin^{d_\xi-2}(t) dt}{\int_0^\pi \sin^{d_\xi-2}(t) dt}.$$

For $0 \leq t \leq \pi/2$, we have $\sin t \geq \frac{2}{\pi}t$. Hence, for $0 < \theta \leq \pi/2$,

$$\begin{aligned} \nu(B_S(I_x^*, \theta)) &\geq \frac{1}{\int_0^\pi \sin^{d_\xi-2}(t) dt} \int_0^\theta \left(\frac{2t}{\pi}\right)^{d_\xi-2} dt \\ &= \frac{(2/\pi)^{d_\xi-2}}{(d_\xi-1) \int_0^\pi \sin^{d_\xi-2}(t) dt} \theta^{d_\xi-1}. \end{aligned}$$

This proves the claim. $\square$

We now optimize the choice of $\theta \in (0, \pi/2]$ to maximize the lower bound $c_{d_\xi} \theta^{d_\xi-1} \cdot \exp\{-L'_\xi \|A\|_{2 \rightarrow 2} \theta/w\}$. Note that this maximum is attained at $\theta = \frac{(d_\xi-1)w}{L'_\xi \|A\|_{2 \rightarrow 2}}$, whenever $w \leq \frac{\pi L'_\xi \|A\|_{2 \rightarrow 2}}{2(d_\xi-1)}$. Therefore, the largest lower bound for $\|dp_x^w/d\nu\|_\infty^{-1}$ is

$$\frac{\left(\frac{2}{\pi}(d_\xi-1)\right)^{d_\xi-2} \exp\{-(d_\xi-1)\}}{\int_0^\pi \sin^{d_\xi-2}(t) dt} \cdot \left(\frac{w}{L'_\xi \|A\|_{2 \rightarrow 2}}\right)^{d_\xi-1} =: C_{d_\xi} \left(\frac{w}{L'_\xi \|A\|_{2 \rightarrow 2}}\right)^{d_\xi-1}. \quad \square$$

EC.4.9. Proof of Theorem 2

We first prove a lemma that characterizes the convergence of the law of I_n to the invariant distribution p_x^w .

LEMMA EC.9. *Under Assumptions 6 and 7, consider the incumbent sequence $\{x_n : n \in \mathbb{N}_+\}$ on a measurable space $(\mathcal{X}, \mathcal{B}_\mathcal{X})$ with respect to the natural filtration $\{\mathcal{F}_n : n \in \mathbb{N}_+\}$ generated by $(\mathcal{I}_m, X_m)_{m \leq n}$:*

$$X_0, I_0 \rightarrow I_1 \rightarrow X_1 \rightarrow I_2 \rightarrow X_2 \rightarrow \dots$$

Conditional on $\mathcal{F}_n$, the next state $\mathcal{I}_{n+1}$ is sampled from $\mathcal{I}_n$ through the one-step Markov kernel P_{x_n} , which admits an invariant distribution $p_{x_n}^w$ on $\mathcal{I}$.

Fix $p > 0$ and suppose that the incumbent movement satisfies $\|X_n - X_{n-1}\| \leq c\alpha_n$ almost surely, where $\{\alpha_n : n \in \mathbb{N}_+\}$ satisfies $\alpha_n \leq C_\alpha n^{-p}$ for some $C_\alpha > 0$. Then, for every sufficiently large n such that $f(n) := \lceil (p+1) \log n / |\log \bar{\eta}| \rceil < n/2$, we have

$$\left\| \mathcal{L}(I_n \mid \mathcal{F}_{n-f(n)}) - p_{X_{n-f(n)}}^w \right\|_{\text{TV}} \leq \bar{\eta}^{f(n)} + \left(\frac{2^p c L_P C_\alpha}{1 - \bar{\eta}} \right) f(n) n^{-p} = O\left(\frac{L_P}{1 - \bar{\eta}} n^{-p} \log n \right).$$

Proof of Lemma EC.9: Fix a sufficiently large n and for $k \geq n - f(n)$, we write $\mu_k := \mathcal{L}(I_k | \mathcal{F}_{n-f(n)})$. Insert the law of a chain that starts from the state $I_{n-f(n)}$ and then applies the frozen kernel $P_{X_{n-f(n)}}$ for $f(n)$ steps. The triangle inequality gives

$$\begin{aligned} & \left\| \mathcal{L}(I_n | \mathcal{F}_{n-f(n)}) - p_{X_{n-f(n)}}^w \right\|_{\text{TV}} \\ & \leq \underbrace{\left\| \delta_{I_{n-f(n)}} P_{X_{n-f(n)}}^{f(n)} - p_{X_{n-f(n)}}^w \right\|_{\text{TV}}}_{\text{(A) mixing error of the frozen chain}} + \underbrace{\left\| \mathcal{L}(I_n | \mathcal{F}_{n-f(n)}) - \delta_{I_{n-f(n)}} P_{X_{n-f(n)}}^{f(n)} \right\|_{\text{TV}}}_{\text{(B) actual endogenous chain versus frozen chain}} \end{aligned}$$

The term (A) measures the mixing error and can be bounded as follows. By definition of the invariant distribution, we have $p_{X_{n-f(n)}}^w P_{X_{n-f(n)}} = p_{X_{n-f(n)}}^w$, where $P_{X_{n-f(n)}}$ is understood as a transition Markov operator defined by $(\mu P_{X_{n-f(n)}})(A) = \int_{\mathcal{I}} P_{X_{n-f(n)}}(I, A) \mu(dI)$ for any measure μ . Then, applying the Dobrushin contraction inequality from Assumption 7 for $f(n)$ times yields

$$\left\| \delta_{I_{n-f(n)}} P_{X_{n-f(n)}}^{f(n)} - p_{X_{n-f(n)}}^w \right\|_{\text{TV}} \leq \bar{\eta}^{f(n)} \left\| \delta_{I_{n-f(n)}} - p_{X_{n-f(n)}}^w \right\|_{\text{TV}} \leq \bar{\eta}^{f(n)}.$$

The term (B) measures the perturbation caused by using the changing kernels $P_{X_{n-f(n)}}, \dots, P_{X_{n-1}}$ instead of repeatedly using the frozen kernel $P_{X_{n-f(n)}}$; its control requires a one-step comparison followed by an induction.

(I) one-step freezing error. For every $k \geq n - f(n)$, we claim that

$$\left\| \mu_{k+1} - \mu_k P_{X_{n-f(n)}} \right\|_{\text{TV}} \leq cL_P \sum_{j=n-f(n)+1}^k \alpha_j.$$

Proof of Claim (I): Fix a Borel set $B \subseteq \mathcal{I}$. Because $\mathcal{F}_{n-f(n)} \subseteq \mathcal{F}_k$, the tower property gives

$$\begin{aligned} \mu_{k+1}(B) &= \mathbb{E}[\mathbf{1}\{I_{k+1} \in B\} | \mathcal{F}_{n-f(n)}] \\ &= \mathbb{E}[\mathbb{E}[\mathbf{1}\{I_{k+1} \in B\} | \mathcal{F}_k] | \mathcal{F}_{n-f(n)}] \\ &= \mathbb{E}[P_{X_k}(I_k, B) | \mathcal{F}_{n-f(n)}]. \end{aligned}$$

Applying the frozen kernel $P_{X_{n-f(n)}}$ to μ_k gives $(\mu_k P_{X_{n-f(n)}})(B) = \mathbb{E}[P_{X_{n-f(n)}}(I_k, B) | \mathcal{F}_{n-f(n)}]$. Subtracting this from the previous representation gives

$$\begin{aligned} \left| \mu_{k+1}(B) - (\mu_k P_{X_{n-f(n)}})(B) \right| &= \left| \mathbb{E}[P_{X_k}(I_k, B) - P_{X_{n-f(n)}}(I_k, B) | \mathcal{F}_{n-f(n)}] \right| \\ &\leq \mathbb{E} \left[\left| P_{X_k}(I_k, B) - P_{X_{n-f(n)}}(I_k, B) \right| | \mathcal{F}_{n-f(n)} \right] \\ &\leq \mathbb{E} \left[\left\| P_{X_k}(I_k, \cdot) - P_{X_{n-f(n)}}(I_k, \cdot) \right\|_{\text{TV}} | \mathcal{F}_{n-f(n)} \right], \end{aligned}$$

where the first inequality follows from conditional Jensen's inequality and taking the supremum over all Borel sets B yields the second inequality. Taking the supremum over all Borel sets B yields

$$\left\| \mu_{k+1} - \mu_k P_{X_{n-f(n)}} \right\|_{\text{TV}} \leq \mathbb{E} \left[\left\| P_{X_k}(I_k, \cdot) - P_{X_{n-f(n)}}(I_k, \cdot) \right\|_{\text{TV}} | \mathcal{F}_{n-f(n)} \right].$$

Now we use the Lipschitz property with constant L_P from Assumption 6 to obtain

$$\left\| \mu_{k+1} - \mu_k P_{X_{n-f(n)}} \right\|_{\text{TV}} \leq L_P \mathbb{E}[\|X_k - X_{n-f(n)}\| \mid \mathcal{F}_{n-f(n)}].$$

Last, we can bound $\|X_k - X_{n-f(n)}\|$ by $\|X_k - X_{n-f(n)}\| \leq \sum_{j=n-f(n)+1}^k \|X_j - X_{j-1}\| \leq c \sum_{j=n-f(n)+1}^k \alpha_j$. Substituting it back to the previous line proves the claim. $\square$

(II) Propagation of the one-step errors by induction. We next claim that, for every integer m with $n - f(n) \leq m \leq n$, it holds that

$$\left\| \mu_m - \delta_{I_{n-f(n)}} P_{X_{n-f(n)}}^{m-(n-f(n))} \right\|_{\text{TV}} \leq cL_P \sum_{k=n-f(n)}^{m-1} \bar{\eta}^{m-1-k} \sum_{j=n-f(n)+1}^k \alpha_j.$$

Proof of Claim (II): We prove this by **induction** on m . For $m = n - f(n)$, the left hand side is zero as $\mu_{n-f(n)} = \delta_{I_{n-f(n)}}$. The right-hand side is non-negative so the *base claim* holds.

Induction step. Now suppose that the claim holds for some $m \in \{n - f(n), \dots, n - 1\}$ and we can write by the triangle inequality

$$\left\| \mu_{m+1} - \delta_{I_{n-f(n)}} P_{X_{n-f(n)}}^{m+1-(n-f(n))} \right\|_{\text{TV}} \leq \left\| \mu_{m+1} - \mu_m P_{X_{n-f(n)}} \right\|_{\text{TV}} + \left\| \mu_m P_{X_{n-f(n)}} - \delta_{I_{n-f(n)}} P_{X_{n-f(n)}}^{m+1-(n-f(n))} \right\|_{\text{TV}}.$$

The first term of RHS can be bounded by Claim (I) above and the second term can be bounded by applying the Dobrushin contraction once under the common frozen kernel $P_{X_{n-f(n)}}$. Then,

$$\left\| \mu_{m+1} - \delta_{I_{n-f(n)}} P_{X_{n-f(n)}}^{m+1-(n-f(n))} \right\|_{\text{TV}} \leq cL_P \sum_{j=n-f(n)+1}^m \alpha_j + \bar{\eta} \left\| \mu_m - \delta_{I_{n-f(n)}} P_{X_{n-f(n)}}^{m-(n-f(n))} \right\|_{\text{TV}}.$$

Now by the induction assumption, this can be further bounded by

$$\begin{aligned} \left\| \mu_{m+1} - \delta_{I_{n-f(n)}} P_{X_{n-f(n)}}^{m+1-(n-f(n))} \right\|_{\text{TV}} &\leq cL_P \sum_{j=n-f(n)+1}^m \alpha_j + cL_P \bar{\eta} \sum_{k=n-f(n)}^{m-1} \bar{\eta}^{m-1-k} \sum_{j=n-f(n)+1}^k \alpha_j \\ &= cL_P \sum_{k=n-f(n)}^m \bar{\eta}^{m-k} \sum_{j=n-f(n)+1}^k \alpha_j. \end{aligned}$$

This is exactly the claim with $m + 1$ in place of m , so the induction step is complete. $\square$

(III) Taking $m = n$ in Claim (II) gives

$$\left\| \mathcal{L}(I_n \mid \mathcal{F}_{n-f(n)}) - \delta_{I_{n-f(n)}} P_{X_{n-f(n)}}^{f(n)} \right\|_{\text{TV}} \leq cL_P \sum_{k=n-f(n)}^{n-1} \bar{\eta}^{n-1-k} \sum_{j=n-f(n)+1}^k \alpha_j \leq \frac{cL_P}{1-\bar{\eta}} \sum_{j=n-f(n)+1}^{n-1} \alpha_j.$$

Combining (I)–(III) produces

$$\left\| \mathcal{L}(I_n \mid \mathcal{F}_{n-f(n)}) - p_{X_{n-f(n)}}^w \right\|_{\text{TV}} \leq \bar{\eta}^{f(n)} + \frac{cL_P}{1-\bar{\eta}} \sum_{j=n-f(n)+1}^{n-1} \alpha_j.$$

Finally, we recall that $f(n)$ is $O(\log n)$ and $f(n) < n/2$ for a sufficiently large n . Therefore, $\sum_{j=n-f(n)+1}^{n-1} \alpha_j \leq C_\alpha f(n) (n/2)^{-p}$. This completes the proof. $\square$

Now we are ready to prove Theorem 2.

Proof of Theorem 2: We first compare the set of adversarial scenarios with respect to $X_{n-f(n)}$ and X_{n-1} . Specifically, we claim that

$$\widehat{I}_{(\tilde{\Delta}-2\varepsilon)/2}(X_{n-f(n)}) \subseteq \widehat{I}_{\tilde{\Delta}-2\varepsilon}(X_{n-1}) \subseteq I_{\tilde{\Delta}}(X_{n-1}). \quad (\text{Inclusion})$$

Proof of the claim: For the first inclusion, we take any $I \in \widehat{I}_{(\tilde{\Delta}-2\varepsilon)/2}(X_{n-f(n)})$ and notice that

$$\begin{aligned} \widehat{Q}_{X_{n-1}}(I) &\geq \widehat{Q}_{X_{n-f(n)}}(I) - L_{Q'}\|X_{n-1} - X_{n-f(n)}\| \\ &\geq \underbrace{\widehat{Q}^*(X_{n-f(n)}) - \frac{\tilde{\Delta} - 2\varepsilon}{2}}_{\text{from } I \in \widehat{I}_{(\tilde{\Delta}-2\varepsilon)/2}(X_{n-f(n)})} - L_{Q'}\|X_{n-1} - X_{n-f(n)}\| \\ &\geq \widehat{Q}^*(X_{n-1}) - \frac{\tilde{\Delta} - 2\varepsilon}{2} - 2L_{Q'}\|X_{n-1} - X_{n-f(n)}\|, \end{aligned}$$

where the last inequality follows from $|\widehat{Q}^*(x) - \widehat{Q}^*(y)| \leq L_{Q'}\|x - y\|$. Indeed, Assumption 5 implies that, for all $I \in \mathcal{I}$, $\widehat{Q}_x(I) \leq \widehat{Q}_y(I) + L_{Q'}\|x - y\| \leq \widehat{Q}^*(y) + L_{Q'}\|x - y\|$. Taking the supremum over I gives one direction of this inequality, and interchanging x and y gives the other direction. Now recall that we have proved in Lemma EC.9 that $\|X_{n-1} - X_{n-f(n)}\| \leq 2^p c C_\alpha f(n) n^{-p}$. Because n is sufficiently large with $2^p c C_\alpha f(n) n^{-p} \leq (\tilde{\Delta} - 2\varepsilon)/(4L'_{Q'})$, we have $\widehat{Q}_{X_{n-1}}(I) \geq \widehat{Q}^*(X_{n-1}) - (\tilde{\Delta} - 2\varepsilon)$ and so $I \in \widehat{I}_{\tilde{\Delta}-2\varepsilon}(X_{n-1})$.

For the second inclusion, we recall that $\widehat{Q}^*(x) \geq Q^*(x) - \varepsilon$. Therefore, if $I \in \widehat{I}_{\tilde{\Delta}-2\varepsilon}(x)$ then $Q(x, I) \geq \widehat{Q}_x(I) - \varepsilon \geq \widehat{Q}^*(x) - (\tilde{\Delta} - 2\varepsilon) - \varepsilon \geq Q^*(x) - \tilde{\Delta}$, implying that $I \in I_{\tilde{\Delta}}(X_{n-1})$. $\square$

Next, we apply the set inclusion claim to bound the target probability. The first inclusion in (Inclusion) implies that $\{I_n \notin I_{\tilde{\Delta}}(X_{n-1})\} \subseteq \{I_n \notin \widehat{I}_{(\tilde{\Delta}-2\varepsilon)/2}(X_{n-f(n)})\}$. Taking conditional probabilities given filtration $\mathcal{F}_{n-f(n)}$ gives

$$\begin{aligned} \mathbb{P}\{I_n \notin I_{\tilde{\Delta}}(X_{n-1}) \mid \mathcal{F}_{n-f(n)}\} &\leq \mathbb{P}\{I_n \notin \widehat{I}_{(\tilde{\Delta}-2\varepsilon)/2}(X_{n-f(n)}) \mid \mathcal{F}_{n-f(n)}\} \\ &= \mathcal{L}(I_n \mid \mathcal{F}_{n-f(n)}) \left(\widehat{I}_{(\tilde{\Delta}-2\varepsilon)/2}(X_{n-f(n)})^c \right). \end{aligned}$$

Recall from Lemma EC.9 that $\|\mathcal{L}(I_n \mid \mathcal{F}_{n-f(n)}) - p_{X_{n-f(n)}}^w\|_{\text{TV}} \leq \bar{\eta}^{f(n)} + ((2^p c L_P C_\alpha)/(1 - \bar{\eta}))f(n)n^{-p}$. Using $\mu(B) \leq \pi(B) + \|\mu - \pi\|_{\text{TV}}$ with $B = \widehat{I}_{(\tilde{\Delta}-2\varepsilon)/2}(X_{n-f(n)})^c$, we obtain

$$\mathbb{P}(I_n \notin I_{\tilde{\Delta}}(X_{n-1}) \mid \mathcal{F}_{n-f(n)}) \leq p_{X_{n-f(n)}}^w \left(\widehat{I}_{(\tilde{\Delta}-2\varepsilon)/2}(X_{n-f(n)})^c \right) + \bar{\eta}^{f(n)} + \frac{2^p c L_P C_\alpha}{1 - \bar{\eta}} f(n) n^{-p}.$$

To bound the first term on the right-hand side, we observe that

$$p_{X_{n-f(n)}}^w \left(\widehat{I}_{(\tilde{\Delta}-2\varepsilon)/2}(X_{n-f(n)})^c \right) = \frac{\int_{\widehat{I}_{(\tilde{\Delta}-2\varepsilon)/2}(X_{n-f(n)})^c} \exp\{\widehat{Q}_{X_{n-f(n)}}(I)/w\} d\nu(I)}{\int_{\mathcal{I}} \exp\{\widehat{Q}_{X_{n-f(n)}}(J)/w\} d\nu(J)}.$$

Now

$$\begin{aligned} \int_{\widehat{I}_{(\tilde{\Delta}-2\varepsilon)/2}(X_{n-f(n)})^c} \exp\left\{\frac{\widehat{Q}_{X_{n-f(n)}}(I)}{w}\right\} d\nu(I) &\leq \nu\left(\widehat{I}_{(\tilde{\Delta}-2\varepsilon)/2}(X_{n-f(n)})^c\right) \sup_{I \notin \widehat{I}_{(\tilde{\Delta}-2\varepsilon)/2}(x)} \exp\left\{\frac{\widehat{Q}_x(I)}{w}\right\} \\ &\leq \nu\left(\widehat{I}_{(\tilde{\Delta}-2\varepsilon)/2}(X_{n-f(n)})^c\right) \exp\left\{\frac{\widehat{Q}^*(X_{n-f(n)}) - (\tilde{\Delta}-2\varepsilon)/2}{w}\right\} \end{aligned}$$

and

$$\begin{aligned} \int_{\mathcal{I}} \exp\left\{\frac{\widehat{Q}_{X_{n-f(n)}}(J)}{w}\right\} d\nu(J) &\geq \nu\left(\widehat{I}_{(\tilde{\Delta}-2\varepsilon)/4}(X_{n-f(n)})\right) \inf_{I \in \widehat{I}_{(\tilde{\Delta}-2\varepsilon)/4}(x)} \exp\left\{\frac{\widehat{Q}_x(I)}{w}\right\} \\ &\geq \nu\left(\widehat{I}_{(\tilde{\Delta}-2\varepsilon)/4}(X_{n-f(n)})\right) \exp\left\{\frac{\widehat{Q}^*(X_{n-f(n)}) - (\tilde{\Delta}-2\varepsilon)/4}{w}\right\}. \end{aligned}$$

Thus,

$$p_{X_{n-f(n)}}^w\left(\widehat{I}_{(\tilde{\Delta}-2\varepsilon)/2}(X_{n-f(n)})^c\right) \leq \frac{\nu(\widehat{I}_{(\tilde{\Delta}-2\varepsilon)/2}(X_{n-f(n)})^c)}{\nu(\widehat{I}_{(\tilde{\Delta}-2\varepsilon)/4}(X_{n-f(n)})^c)} \exp\left\{-\frac{\tilde{\Delta}-2\varepsilon}{4w}\right\}.$$

Last, we summarize and remove the conditioning to obtain

$$\begin{aligned} \mathbb{P}(I_n \notin I_{\tilde{\Delta}}(X_{n-1})) &= \mathbb{E}[\mathbb{P}(I_n \notin I_{\tilde{\Delta}}(X_{n-1}) \mid \mathcal{F}_{n-f(n)})] \\ &\leq \bar{\eta}^{f(n)} + \frac{2^p c L_P C_\alpha}{1 - \bar{\eta}} f(n) n^{-p} + \sup_{x \in \mathcal{X}} \frac{\nu(\widehat{I}_{(\tilde{\Delta}-2\varepsilon)/2}(x)^c)}{\nu(\widehat{I}_{(\tilde{\Delta}-2\varepsilon)/4}(x))} \exp\left\{-\frac{\tilde{\Delta}-2\varepsilon}{4w}\right\}. \quad \square \end{aligned}$$

EC.4.10. Proof of Example 10

Proof: Recall from Lemma 1 that $Q(x, \cdot)$ is L'_ξ -Lipschitz on Ξ , uniformly over $x \in \mathcal{X}$. Since $\widehat{Q}_x(I) = Q(x, \xi_I)$ and $\xi_I = I$ in Example 10, it follows that $\widehat{Q}_x(\cdot)$ is also L'_ξ -Lipschitz on $\mathcal{I} = \Xi$.

We first establish a uniform lower bound on the measure of metric balls. Equip Ξ with the Euclidean metric and recall that ν is the normalized Lebesgue measure on Ξ , i.e., $\nu(A) := \frac{\text{vol}_{d_\xi}(A \cap \Xi)}{\text{vol}_{d_\xi}(\Xi)}$. For any $I \in \Xi$ and $0 < r \leq D_{\mathcal{I}}$, define the scaled set $S := I + (r/D_{\mathcal{I}})(\Xi - I)$. By convexity of Ξ , $S \subseteq \Xi$. Moreover, for any $s \in S$, there exists $J \in \Xi$ such that $s = I + (r/D_{\mathcal{I}})(J - I)$, and hence $\|s - I\| = \frac{r}{D_{\mathcal{I}}} \|J - I\| \leq \frac{r}{D_{\mathcal{I}}} \text{diam}(\Xi) = r$. Thus, $S \subseteq B_{\mathcal{I}}(I, r) \cap \Xi$, where $B_{\mathcal{I}}(I, r) := \{J \in \mathcal{I} : \|J - I\| \leq r\}$. By the transformation properties of the Lebesgue measure under scaling and translation (see [Rudin 1987](#), Theorem 2.20, 2.23), $\text{vol}_{d_\xi}(S) = (r/D_{\mathcal{I}})^{d_\xi} \text{vol}_{d_\xi}(\Xi)$. Consequently,

$$\nu(B_{\mathcal{I}}(I, r)) \geq \frac{\text{vol}_{d_\xi}(S)}{\text{vol}_{d_\xi}(\Xi)} = \left(\frac{r}{D_{\mathcal{I}}}\right)^{d_\xi}.$$

For any $x \in \mathcal{X}$, let $I_x^* \in \arg \max_{I \in \mathcal{I}} \widehat{Q}_x(I)$ and set $r_{\tilde{\Delta}} := \min\{D_{\mathcal{I}}, \tilde{\Delta}/(4L'_\xi)\}$. Since $L'_\xi r_{\tilde{\Delta}} \leq \tilde{\Delta}/4$, the L'_ξ -Lipschitz continuity of $\widehat{Q}_x(\cdot)$ implies $B_{\mathcal{I}}(I_x^*, r_{\tilde{\Delta}}) \subseteq \widehat{I}_{\tilde{\Delta}/4}(x)$. Therefore, by the preceding small-ball bound and $\nu(\mathcal{I}) = 1$,

$$\nu(\widehat{I}_{\tilde{\Delta}/4}(x)) \geq \left(\frac{r_{\tilde{\Delta}}}{D_{\mathcal{I}}}\right)^{d_\xi} = \min\left\{1, \frac{\tilde{\Delta}}{4L'_\xi D_{\mathcal{I}}}\right\}^{d_\xi}, \implies \frac{\nu(\widehat{I}_{\tilde{\Delta}/2}(x)^c)}{\nu(\widehat{I}_{\tilde{\Delta}/4}(x))} \leq \max\left\{1, \frac{4L'_\xi D_{\mathcal{I}}}{\tilde{\Delta}}\right\}^{d_\xi}.$$

Taking the supremum over $x \in \mathcal{X}$ completes the proof.

EC.4.11. Proof of Theorem 3

Proof: We denote $I_t^* := I^*(X_t)$ for $t \in \mathbb{N}_+$. We first bound $\mathbb{P}(I_{t+1} \notin I_t^* \mid \mathcal{F}_t)$ by discussing the following two cases.

Case 1. If $I_t \in I_t^*$, then for any $j \notin I_t^*$,

$$\widehat{Q}_{X_t}(I_t) - \widehat{Q}_{X_t}(j) \geq Q(X_t, I_t) - Q(X_t, j) - 2\varepsilon \geq \tilde{\Delta}_t - 2\varepsilon \geq \Delta_* - 2\varepsilon.$$

The invariance of $p_{X_t}^w$ gives $p_{X_t}^w((I_t^*)^c) = \sum_{i \in \mathcal{I}} p_{X_t}^w(i) P_{X_t}(i, (I_t^*)^c) \geq p_{X_t}^w(I_t) P_{X_t}(I_t, (I_t^*)^c)$. Then,

$$\mathbb{P}(I_{t+1} \notin I_t^* \mid \mathcal{F}_t) \leq \frac{p_{X_t}^w((I_t^*)^c)}{p_{X_t}^w(I_t)} = \sum_{j \notin I_t^*} \exp \left\{ \frac{\widehat{Q}_{X_t}(j) - \widehat{Q}_{X_t}(I_t)}{w} \right\} \leq (|\mathcal{I}| - 1) \exp \left\{ -\frac{\Delta_* - 2\varepsilon}{w} \right\}.$$

Case 2. If $I_t \notin I_t^*$, then we choose any $j \in I_t^*$ and it follows from the definition of total variation that

$$\begin{aligned} \mathbb{P}(I_{t+1} \notin I_t^* \mid \mathcal{F}_t) &= P_{X_t}(I_t, (I_t^*)^c) \leq P_{X_t}(j, (I_t^*)^c) + \|P_{X_t}(I_t, \cdot) - P_{X_t}(j, \cdot)\|_{\text{TV}} \\ &\leq (|\mathcal{I}| - 1) \exp \left\{ -\frac{\Delta_* - 2\varepsilon}{w} \right\} + \bar{\eta}, \end{aligned}$$

where the second inequality follows from Case 1 and Assumption 7. Combining these two cases gives

$$\mathbb{P}(I_{t+1} \notin I_t^* \mid \mathcal{F}_t) \leq (|\mathcal{I}| - 1) \exp \left\{ -\frac{\Delta_* - 2\varepsilon}{w} \right\} + \bar{\eta} \mathbf{1}_{\{I_t \notin I_t^*\}}.$$

We then take expectations and use the tower property to obtain $\mathbb{P}(I_{t+1} \notin I_t^*) \leq (|\mathcal{I}| - 1) \exp \left\{ -(\Delta_* - 2\varepsilon)/w \right\} + \bar{\eta} \mathbb{P}\{I_t \notin I_t^*\}$. If $I_t \notin I_t^*$, then either $I_t \notin I_{t-1}^*$ or $I_t^* \neq I_{t-1}^*$ (if neither event happens, then $I_t \in I_{t-1}^* = I_t^*$, contradicting the presumption). It follows that

$$\mathbb{P}\{I_t \notin I_t^*\} \leq \mathbb{P}\{I_t \notin I_{t-1}^*\} + \mathbb{P}\{I_t^* \neq I_{t-1}^*\}.$$

Second, we bound the second term in right-hand side above. If two sets I_t^* and I_{t-1}^* differ, then WLOG there is a state $a \in I_{t-1}^* \setminus I_t^*$ (if $I_{t-1}^* \setminus I_t^* = \emptyset$, then pick an $a \in I_t^* \setminus I_{t-1}^*$ and a similar proof as below works). For any state $b \in I_t^*$, we have

$$\widehat{Q}_{X_t}(b) - \widehat{Q}_{X_t}(a) \geq \Delta_* - 2\varepsilon, \quad \widehat{Q}_{X_{t-1}}(b) - \widehat{Q}_{X_{t-1}}(a) \leq 2\varepsilon.$$

It follows that

$$\begin{aligned} \Delta_* - 4\varepsilon &\leq [\widehat{Q}_{X_t}(b) - \widehat{Q}_{X_t}(a)] - [\widehat{Q}_{X_{t-1}}(b) - \widehat{Q}_{X_{t-1}}(a)] \\ &= [\widehat{Q}_{X_t}(b) - \widehat{Q}_{X_{t-1}}(b)] + [\widehat{Q}_{X_{t-1}}(a) - \widehat{Q}_{X_t}(a)] \\ &\leq \left| \widehat{Q}_{X_{t-1}}(b) - \widehat{Q}_{X_t}(b) \right| + \left| \widehat{Q}_{X_t}(a) - \widehat{Q}_{X_{t-1}}(a) \right|. \end{aligned}$$

Thus,

$$\begin{aligned} \mathbb{P}\{I_t^* \neq I_{t-1}^*\} &= \mathbb{E}[\mathbf{1}\{I_t^* \neq I_{t-1}^*\}] \leq \mathbb{E}\left[\frac{\left|\widehat{Q}_{X_{t-1}}(b) - \widehat{Q}_{X_t}(b)\right| + \left|\widehat{Q}_{X_t}(a) - \widehat{Q}_{X_{t-1}}(a)\right|}{\Delta_* - 4\varepsilon}\right] \\ &\leq \frac{2L_{Q'}\|X_t - X_{t-1}\|}{\Delta_* - 4\varepsilon} \leq \frac{2cL'_Q}{\Delta_* - 4\varepsilon}\alpha_t. \end{aligned}$$

We combine the results and obtain

$$\mathbb{P}(I_{t+1} \notin I_t^*) \leq (|\mathcal{I}| - 1) \exp\left\{-\frac{\Delta_* - 2\varepsilon}{w}\right\} + \bar{\eta} \left(\mathbb{P}\{I_t \notin I_{t-1}^*\} + \frac{2cL'_Q}{\Delta_* - 4\varepsilon}\alpha_t\right).$$

Telescoping this inequality for $t = 1$ through $t = n - 1$ gives a geometric coefficient $\bar{\eta}^{n-1-t}$ for the movement term generated at time t , and a geometric sum $\sum_{j=0}^{n-2} \bar{\eta}^j = (1 - \bar{\eta}^{n-1})/(1 - \bar{\eta})$. In other words, for any $n \geq 2$, we obtain

$$\mathbb{P}(I_n \notin I^*(X_{n-1})) \leq \bar{\eta}^{n-1} \mathbb{P}(I_1 \notin I^*(X_0)) + \frac{2c\bar{\eta}L'_Q}{\Delta_* - 4\varepsilon} \sum_{t=1}^{n-1} \bar{\eta}^{n-1-t} \alpha_t + \frac{|\mathcal{I}| - 1}{1 - \bar{\eta}} (1 - \bar{\eta}^{n-1}) \exp\left\{-\frac{\Delta_* - 2\varepsilon}{w}\right\},$$

which proves the first claim of Theorem 3.

Finally, recall that the preceding argument on the event $\{I_t^* \neq I_{t-1}^*\}$ implies

$$\Delta_* - 4\varepsilon \leq 2L'_Q\|X_t - X_{t-1}\| \leq 2cC_\alpha L'_Q t^{-p}.$$

But, for all $t \geq n_0 := \max\left\{2, 1 + \left\lceil \left(\frac{2cC_\alpha L'_Q}{\Delta_* - 4\varepsilon}\right)^{1/p} \right\rceil\right\}$, the last quantity is strictly smaller than $\Delta_* - 4\varepsilon$. Therefore, $\mathbb{P}(I_t^* \neq I_{t-1}^*) = 0$ for all $t \geq n_0$. The prior recursion therefore simplifies to

$$\mathbb{P}(I_{t+1} \notin I_t^*) \leq \bar{\eta} \mathbb{P}(I_t \notin I_{t-1}^*) + (|\mathcal{I}| - 1) \exp\left\{-\frac{\Delta_* - 2\varepsilon}{w}\right\}.$$

Telescoping this inequality, starting from $t = n_0$, produces the second claim of Theorem 3. $\square$

EC.4.12. Proof of Theorem 4

Proof: We shall keep using the framework for proving Theorem 1. It suffices to justify its extension from finite matrices to operators on bounded measurable functions and verify the corresponding constants. The major descent and averaging arguments remain similar. We note that the score is exact, i.e., $\widehat{Q}_x(I) \equiv Q(x, \xi_I)$ (and hence $\epsilon = 0$), and let ν be a probability measure.

(i) Properties inherited from Lemmas 1 and 3.

First, the proof of Lemma 1 applies uniformly to all $I \in \mathcal{I}$ as the dual polyhedron is independent of the state space. It follows that $-E^\top \pi^*(x, I) \in \partial_x Q(x, \xi_I)$, $\|E^\top \pi^*(x, I)\| \leq L'_Q$. Since ν is a finite nonzero measure and $Q(x, \cdot)$ is measurable and bounded for every $x \in \mathcal{X}$, we have $0 < \int_{\mathcal{I}} \exp\{Q(x, \xi_I)/w\} \nu(dI) < \infty$. Hence, F_w^ν and p_x^w are well-defined.

Second, the proof of Lemma 3 extends by replacing finite sums with integrals (note that the selected subgradients are measurable and uniformly bounded, so their integrals are well-defined). In particular, we have $z_w(x) := \int_{\mathcal{I}} Z_x(I) p_x^w(dI) \in \partial F_w^\nu(x)$, $\|Z_x(I)\|, \|z_w(x)\| \leq L_f + L'_Q$, and $|F_w^\nu(x) - F_w^\nu(y)| \leq (L_f + L'_Q)\|x - y\|$ for all $x, y \in \mathcal{X}$.

(ii) The continuous-state Poisson operator and Proposition 1.

For a bounded measurable function φ , define $(P_x \varphi)(I) := \int \varphi(J) P_x(I, dJ)$, and $((\mathbf{1} p_x^w) \varphi)(I) := \int \varphi(J) p_x^w(dJ)$. Note that these are bounded operators under the supremum norm and the MH kernel has invariant measure p_x^w by the detailed balance equations. For any probability measures μ and λ on $\mathcal{I}$ and for all $t \geq 1$, Assumption 7 implies that

$$\|\mu P_x^t - \lambda P_x^t\|_{\text{TV}} \leq \bar{\eta}^t \|\mu - \lambda\|_{\text{TV}}.$$

Since p_x^w is invariant for P_x , taking $\mu = \delta_I$ and $\lambda = p_x^w$ yields

$$\|P_x^t(I, \cdot) - p_x^w\|_{\text{TV}} = \|\delta_I P_x^t - p_x^w P_x^t\|_{\text{TV}} \leq \bar{\eta}^t \|\delta_I - p_x^w\|_{\text{TV}} \leq \bar{\eta}^t,$$

where the first equality uses the invariance of p_x^w and the last inequality uses the fact that $\|\mu - \lambda\|_{\text{TV}} := \sup_B |\mu(B) - \lambda(B)| \leq 1$ for probability measures. Recalling the dual representation $\sup_{\|\varphi\|_\infty \leq 1} \left| \int_{\mathcal{I}} \varphi(J) (\mu - \lambda)(dJ) \right| = 2\|\mu - \lambda\|_{\text{TV}}$ of the total variation, we also have

$$\|P_x^t - \mathbf{1} p_x^w\|_{\infty \rightarrow \infty} = \sup_{\|\varphi\|_\infty \leq 1} \sup_{I \in \mathcal{I}} \left| \int_{\mathcal{I}} \varphi(J) (P_x^t(I, dJ) - p_x^w(dJ)) \right| = 2 \sup_{I \in \mathcal{I}} \|P_x^t(I, \cdot) - p_x^w\|_{\text{TV}} \leq 2\bar{\eta}^t.$$

But the operator $(\text{Id} - P_x + \mathbf{1} p_x^w)^{-1}$ admits the equality $(\text{Id} - P_x + \mathbf{1} p_x^w)^{-1} = \text{Id} + \sum_{t=1}^{\infty} (P_x^t - \mathbf{1} p_x^w)$. To see this, multiply the partial sum $\left[\text{Id} + \sum_{t=1}^T (P_x^t - \mathbf{1} p_x^w) \right]$ by $(\text{Id} - P_x + \mathbf{1} p_x^w)$ to obtain $\text{Id} - (P_x^{T+1} - \mathbf{1} p_x^w)$, which converges to Id as T increases. Consequently,

$$\sup_x \|(\text{Id} - P_x + \mathbf{1} p_x^w)^{-1}\|_{\infty \rightarrow \infty} \leq 1 + 2 \sum_{t=1}^{\infty} \bar{\eta}^t = \frac{1 + \bar{\eta}}{1 - \bar{\eta}} := C_*.$$

Note here that C_* carries the same meaning as in Theorem 1, with the finite-dimensional inverse replaced by its bounded-operator counterpart.

Fix $x^* \in \arg \min_{x \in \mathcal{X}} F_w^\nu(x)$ and adapt the centered error in Proposition 1 through

$$e_x(I) := Q(x, \xi_I) - Q(x^*, \xi_I) - \int_{\mathcal{I}} [Q(x, \xi_J) - Q(x^*, \xi_J)] p_x^w(dJ), \quad e_k := e_{x_{k-1}}(I_k).$$

Then, $\int e_x dp_x^w = 0$ and $\|e_x\|_\infty \leq 2DL'_Q$. For each $I \in \mathcal{I}$, define

$$\hat{Z}_x(I) := [(\text{Id} - P_x + \mathbf{1} p_x^w)^{-1} e_x](I) = \sum_{t=0}^{\infty} (P_x^t e_x)(I),$$

which satisfies $\widehat{Z}_x(I) - (P_x \widehat{Z}_x)(I) = e_x(I)$ and $\int \widehat{Z}_x dp_x^w = 0$. We can then similarly bound it as

$$\max \left\{ \sup_I |\widehat{Z}_x(I)|, \sup_I |(P_x \widehat{Z}_x)(I)| \right\} \leq B_Z := 2C_* DL'_Q.$$

Lemma [EC.6](#), which we proved for a general state space, gives that

$$\|\mathbf{1}p_x^w - \mathbf{1}p_y^w\|_{\infty \rightarrow \infty} \leq L_P \|x - y\| \text{ with } L_P := \frac{2L'_Q}{w}$$

and Assumption [6](#) gives $\|P_x - P_y\|_{\infty \rightarrow \infty} \leq 2L_P \|x - y\|$ for all $x, y \in \mathcal{X}$ (note that the factor of 2 follows from the total-variation convention in Assumption [6](#)). Furthermore, exactly as for the centered error in Lemma [EC.7](#), we have $\|e_x - e_y\|_{\infty} \leq (2L'_Q + DL'_Q L_P) \|x - y\|$. The resolvent identity used in the proof of Lemma [EC.7](#) remains valid for bounded invertible operators. Applying that same argument yields

$$\sup_I |\widehat{Z}_x(I) - \widehat{Z}_y(I)| \leq [C_*(2L'_Q + DL'_Q L_P) + 2C_*^2 DL'_Q (2L_P + L_P)] \|x - y\| := L_Z \|x - y\|.$$

Adding and subtracting $P_x \widehat{Z}_y$ then gives

$$\sup_I |(P_x \widehat{Z}_x)(I) - (P_y \widehat{Z}_y)(I)| \leq (2L_P B_Z + L_Z) \|x - y\|.$$

Therefore, the bounds required by Proposition [1](#) still hold, with the only change being its kernel operator modulus L_P replaced by $2L_P$.

(iii) Application of Proposition [1](#) and Theorem [1](#).

The remainder of the proof uses exactly the same arguments as in the proofs of Proposition [1](#) and Theorem [1](#), with F_w replaced by F_w^ν and the constants defined above. They therefore yield the stated bounds for both the step-size average and the linear average. Substituting $\alpha_k = \alpha_0/\sqrt{k}$ and $L_P = 2L'_Q/w$ gives the claimed rates. $\square$

EC.4.13. Proof of Corollary [1](#)

Proof: For the sequence associated with the fixed integer, let X_k be the incumbent solution after k previous oracle calls and let I_k be the state stored immediately before this update. Thus the update at X_k returns I_{k+1} . The initial call, indexed by $k = 0$, returns I_1 and initializes the recursion. For every $k \geq 1$, if $\|X_k - X_{k-1}\| \leq c_0 k^{-p}$, the argument in the proof of Theorem [3](#) gives

$$\mathbb{P}(I_{k+1} \notin I^*(X_k)) \leq \bar{\eta} \mathbb{P}(I_k \notin I^*(X_{k-1})) + \frac{2c_0 \bar{\eta} L_{Q'}}{\Delta_* - 4\varepsilon} k^{-p} + (|I| - 1) \exp \left\{ -\frac{\Delta_* - 2\varepsilon}{w} \right\}.$$

Suppose instead that the movement condition fails. The incumbent remains fixed during the ensuing $g(\kappa)$ transitions. The Dobrushin contraction and invariance of $p_{X_k}^w$ imply, uniformly over the initial state,

$$\left\| \delta_{I_k} P_{X_k}^{g(\kappa)} - p_{X_k}^w \right\|_{\text{TV}} \leq \bar{\eta}^{g(\kappa)}.$$

The uniform score gap and the approximation error remain

$$p_{X_k}^w(I^*(X_k)^c) \leq (|\mathcal{I}| - 1) \exp\left\{-\frac{\Delta_* - 2\varepsilon}{w}\right\}.$$

At the transition indexed by $\kappa \geq 1$, exactly κ previous oracle calls have been completed, so $\max\{\kappa(x_z), 1\} = \kappa$. We choose $g(\kappa)$ such that

$$\bar{\eta}^{g(\kappa)} \leq \frac{2c_0\bar{\eta}L_{Q'}}{\Delta_* - 4\varepsilon} \kappa^{-p}.$$

For our case, we therefore choose

$$g(\kappa(x_z)) := \max\left\{1, \left\lceil \frac{p \log(\max\{\kappa(x_z), 1\}) - \log\left(\frac{2c_0\bar{\eta}L_{Q'}}{\Delta_* - 4\varepsilon}\right)}{|\log \bar{\eta}|} \right\rceil \right\}.$$

Consequently, the one-step recursive argument in the proof of Theorem 3 applies. After κ calls, $\mathbf{state}(\bar{x}_z) = I_\kappa$ and $\mathbf{ref}(\bar{x}_z)$ is the continuous component of $X_{\kappa-1}$. Iterating from $k = 1$ through $k = \kappa - 1$ gives us the same result as in Theorem 3.

It remains to obtain a convenient closed-form bound. For $\kappa \geq 2$, splitting its index set at $\lfloor (\kappa - 1)/2 \rfloor$ gives

$$\sum_{t=1}^{\kappa-1} \bar{\eta}^{\kappa-1-t} t^{-p} \leq \frac{\bar{\eta}^{\lceil (\kappa-1)/2 \rceil}}{1 - \bar{\eta}} + \frac{2^p}{1 - \bar{\eta}} \kappa^{-p}.$$

Moreover,

$$\bar{\eta}^{\lceil (\kappa-1)/2 \rceil} \leq \bar{\eta}^{-1/2} \exp\left\{-\frac{|\log \bar{\eta}|}{2} \kappa\right\} \leq \bar{\eta}^{-1/2} \left(\frac{2p}{e|\log \bar{\eta}|}\right)^p \kappa^{-p},$$

where the last inequality follows by maximizing $u^p \exp\{-|\log \bar{\eta}|u/2\}$ over $u > 0$. Substitution yields

$$\sum_{t=1}^{\kappa-1} \bar{\eta}^{\kappa-1-t} t^{-p} \leq C_{\bar{\eta},p} \kappa^{-p}. \quad \square$$