

A Shrinkage Path Heuristic for Wasserstein Distributionally Robust Optimization

Lingjun Meng, Ryan Cory-Wright, Wolfram Wiesemann
Imperial Business School, Imperial College London, London SW7 2AZ, United Kingdom
{l.meng23, r.cory-wright, ww}@imperial.ac.uk

Wasserstein distributionally robust optimization (DRO) is a versatile and widely adopted framework for decision-making under uncertainty, yet its standard deterministic reformulations generally contain non-convex inner subproblems that are challenging to solve. To address this issue, we propose a shrinkage path heuristic that reduces the solution of a DRO problem to a one-dimensional search over the line segment connecting the (typically benign) sample average approximation (SAA) and the (more demanding but practically solvable) classical robust optimization solution. We derive a priori suboptimality bounds in stylized settings and, for the general case, a posteriori bounds obtained by applying a similar heuristic to a dual formulation. Numerical experiments on a multi-item newsvendor and an appointment scheduling problem show that the shrinkage path heuristic attains 85–110% (resp. 45–70%) of the out-of-sample performance improvements of Wasserstein DRO over SAA, at a fraction of the computational cost.

1. Introduction

Decision problems under uncertainty are often solved via sample average approximation (SAA), which involves optimizing against the empirical distribution of these observations. However, SAA is prone to overfitting, as the resulting decisions are tailored to the sample rather than reflecting the underlying distribution that generated it, often resulting in poor out-of-sample performance. To counteract this overfitting, distributionally robust optimization (DRO) hedges against all distributions in a neighborhood of the empirical distribution. This protection comes at a price, however. Although DRO problems can be reformulated in certain special cases, they are generally more computationally challenging to solve than their SAA counterparts. Additionally, the size of the neighborhood is a hyperparameter that is rarely known in advance and usually requires tuning through cross-validation, which involves resolving the DRO problem for each candidate neighborhood size across multiple folds. This paper proposes a simple shrinkage path heuristic that replaces DRO with a one-dimensional search along the line segment connecting the solutions of two well-studied problems: the SAA problem and its classical robust optimization counterpart. Since the line segment itself does not depend on any hyperparameters, tuning the degree of robustness via cross-validation requires only the two endpoint solutions and inexpensive evaluations along

the path. The heuristic thereby captures much of the protection of DRO over SAA while reducing computational costs by orders of magnitude.

Specifically, we study DRO problems of the form

$$\begin{aligned} & \underset{\mathbf{x} \in \mathcal{X}}{\text{minimize}} && \sup_{\mathbb{Q} \in \mathbb{B}_\rho(\hat{\mathbb{P}})} \mathbb{E}_{\mathbb{Q}}[\ell(\mathbf{x}, \tilde{\boldsymbol{\xi}})] \\ & \text{subject to} && g_m(\mathbf{x}, \tilde{\boldsymbol{\xi}}) \leq \gamma_m \quad \mathbb{Q}\text{-a.s. } \forall \mathbb{Q} \in \mathbb{B}_\rho(\hat{\mathbb{P}}), \forall m \in [M], \end{aligned} \tag{1}$$

where the decision $\mathbf{x}$ is selected from a closed convex set $\mathcal{X} \subseteq \mathbb{R}^n$, the objective minimizes the worst-case expectation of a loss function ℓ that is closed convex in $\mathbf{x}$ for every $\boldsymbol{\xi} \in \Xi$ and upper semicontinuous in $\boldsymbol{\xi}$ for every $\mathbf{x} \in \mathcal{X}$, and the constraints indexed by $m \in [M] := \{1, \dots, M\}$ are closed convex in $\mathbf{x}$ for every fixed realization $\boldsymbol{\xi}$ and hold $\mathbb{Q}$ -almost surely. The uncertainty is captured by the random vector $\tilde{\boldsymbol{\xi}} \in \Xi$, which may be governed by any distribution $\mathbb{Q}$ from the type-1 Wasserstein ambiguity set

$$\mathbb{B}_\rho(\hat{\mathbb{P}}) := \{\mathbb{Q} \in \mathcal{P}(\Xi) : W(\mathbb{Q}, \hat{\mathbb{P}}) \leq \rho\} \quad \text{with} \quad W(\mathbb{Q}, \hat{\mathbb{P}}) = \inf_{\pi \in \Pi(\mathbb{Q}, \hat{\mathbb{P}})} \int_{\Xi \times \Xi} d(\boldsymbol{\xi}, \boldsymbol{\xi}') \pi(d\boldsymbol{\xi}, d\boldsymbol{\xi}').$$

Here, d is a lower semicontinuous ground metric, and the empirical distribution $\hat{\mathbb{P}} = \frac{1}{N} \sum_{i=1}^N \delta_{\boldsymbol{\xi}_i}$ is constructed from N independent samples $\boldsymbol{\xi}_1, \dots, \boldsymbol{\xi}_N$ drawn from the unknown true data-generating distribution $\mathbb{P}^*$. We assume throughout that the support $\Xi \subseteq \mathbb{R}^k$ is non-empty, closed, and convex.

We refer to Problem (1) as the Wasserstein DRO (WDRO) problem. It arises in a wide range of applications, including inventory management (Hanasusanto et al. 2015, Xin and Goldberg 2022), machine learning (Bennouna et al. 2023), energy systems (Xie and Ahmed 2018, Esteban-Pérez and Morales 2023), and control theory (Shafieezadeh-Abadeh et al. 2018). We refer to Rahimian and Mehrotra (2022) and Kuhn et al. (2025) for detailed reviews.

Unfortunately, Problem (1) is often not *practically tractable*—that is, solvable in a practical amount of time at instance sizes relevant to the application (Bertsimas and Dunn 2019). By strong duality, the worst-case expectation objective admits a finite-dimensional dual reformulation. In this reformulation, one optimizes over the dual multiplier corresponding to the Wasserstein radius and solves N auxiliary maximization problems, one for each sample point, over the support Ξ . Each auxiliary problem maximizes a loss-like term minus a transport-cost penalty (Mohajerin Esfahani and Kuhn 2018, Theorem 4.2). The tractability of Problem (1), therefore, hinges on the structure of these inner maximization problems. When Ξ is unbounded (e.g., $\Xi = \mathbb{R}^k$ or $\Xi = \mathbb{R}_+^k$) and the ground metric is a norm, each supremum is finite only if the loss function is Lipschitz continuous on Ξ . In the special case where $\Xi = \mathbb{R}^k$, this finite problem reduces to a regularized sample average approximation (SAA) (Shafieezadeh-Abadeh et al. 2019, Theorem 4(ii)). However, many losses of practical interest—including quadratic, exponential, and other smooth superlinear losses—fail the Lipschitz condition and therefore can render the worst-case expectation infinite along unbounded

rays. When Ξ is compact, the inner suprema are always finite, but their tractability depends on the loss structure. If the loss function is *a pointwise maximum of finitely many concave functions* of ξ (Mohajerin Esfahani and Kuhn 2018, Assumption 4.1), each inner problem maximizes a piecewise concave function over a convex set, which is tractable. In contrast, when the loss function is *convex* in ξ —as is the case for quadratic and other smooth losses commonly encountered in practice—each inner subproblem becomes a k -dimensional difference-of-convex program; in general, solving such problems is NP-hard, even for highly structured quadratic instances (Pardalos and Vavasis 1991).

Several strategies aim to circumvent the intractability of WDRO with general loss functions. When the loss is convex in ξ , Corollary 4.3 of Mohajerin Esfahani and Kuhn (2018) yields finite convex conservative approximations of the inner suprema, although their sizes typically grow exponentially in the dimension of ξ . Cheramin et al. (2022) propose inner and outer approximations for moment and Wasserstein ambiguity sets based on block decompositions of ξ and principal component reduction, together with gap bounds that quantify the resulting accuracy–runtime trade-off. Gao et al. (2024, Theorem 2) connect type-1 Wasserstein DRO to variation regularization and show that the gap between the worst-case expectation and the empirical loss is bounded by ρ times a variation norm of the loss function.

A second stream of work solves Problem (1) algorithmically. Shafiee et al. (2026) study gradient-based methods for Wasserstein DRO when Ξ is unconstrained and the loss ℓ is nonconvex in x . Wang et al. (2026) replace the Wasserstein distance with its entropically regularized counterpart; the resulting Sinkhorn dual is amenable to optimization via stochastic mirror descent. Two further works propose dual decomposition algorithms for special problem classes: Li et al. (2019) propose an ADMM framework that exploits a dual decomposition for the Wasserstein distributionally robust logistic regression. Luo and Mehrotra (2019) develop a cutting-surface decomposition algorithm for the dual reformulation of the Wasserstein DRO with bounded support for a class of statistical learning problems. Finally, Blanchet et al. (2022) study optimal-transport DRO with quadratic transport costs and linear decision rules, which together reduce each inner adversarial subproblem to a line search, and develop stochastic gradient descent schemes for the outer decision with convergence guarantees. The GitHub repository accompanying this paper classifies the literature in greater depth according to the scope of each approach, the nature of the resulting reformulation or algorithm, and whether the approach is exact¹.

All of the above approaches aim to solve Problem (1)—or an approximation thereof—for a *single fixed* Wasserstein radius ρ . In contrast, we exploit the structure of the solution path *as ρ varies*. We adopt this perspective because, in most practical applications, ρ is not selected a priori (*e.g.*, from a concentration inequality), but rather it is tuned via cross-validation on the data. To this end, we propose a shrinkage path heuristic (Algorithm 1) that reduces Problem (1) to a one-dimensional

search over the line segment that connects two benchmark decisions: the solution of the sample average approximation (SAA), which arises at $\rho = 0$ and is typically easy to compute, and the solution of the classical robust optimization (RO) problem, which arises as $\rho \rightarrow \infty$ and—although harder—is routinely solved in practice. The line segment itself is independent of ρ : instead of solving a full DRO problem at every candidate radius, one tunes the position on this fixed path by cross-validation, saving orders of magnitude in runtime. Although simple, the heuristic admits a priori guarantees in stylized settings (Theorems 1–3 and Proposition 1 in Section 3), a posteriori guarantees in general (Theorems 4 and 5 in Section 4), and competitive empirical performance (Section 5); these three advances form the core of this paper.

More specifically, we summarize our contributions as follows.

1. Under suitable stylized assumptions on the loss function and uncertainty support, we establish *a priori suboptimality bounds* that serve as a theoretical motivation for our heuristic.
2. For the general setting of Problem (1), we apply the same shrinkage idea to a dual formulation, yielding a companion procedure (Algorithm 2) that mirrors Algorithm 1 on the dual side. The gap between the primal and dual heuristic values provides computable (per-instance) *a posteriori suboptimality bounds* without requiring an exact solution of Problem (1).
3. We demonstrate the effectiveness of our heuristic on a multi-item newsvendor and an appointment scheduling problem, where our a posteriori bounds certify that the heuristic incurs only a small in-sample suboptimality gap for the DRO objective. Out of sample, for the multi-item newsvendor setting (resp. the appointment scheduling problem), the shrinkage path heuristic attains 85–110% (resp. 45–70%) of the median out-of-sample benefits of Wasserstein DRO over SAA, at a fraction of the computational cost.

The remainder of this paper is organized as follows. Section 2 introduces our heuristic. Section 3 derives a priori suboptimality bounds under several structural assumptions on the loss function. Section 4 develops a posteriori performance guarantees for the general setting by constructing a dual shrinkage path whose objective values provide computable lower bounds on the true DRO optimum. Section 5 illustrates our heuristic on a multi-item newsvendor and an appointment scheduling problem. All proofs are deferred to the electronic companion. We relegate an extended literature review and a generalization of Problem (1) to ambiguous conditional value-at-risk constraints to the GitHub repository accompanying this work.

REMARK 1 (MULTI-STAGE PROBLEMS). Our heuristic extends to two-stage and multi-stage variants of Problem (1) in two natural ways. The first is to restrict the recourse decisions to affine (or piecewise-affine) functions of the uncertainty, which collapses the problem to a single-stage instance of Problem (1) in an expanded first-stage decision vector. The second is to substitute the optimal recourse value function for ℓ directly, whenever this value function remains convex in $(\mathbf{x}, \boldsymbol{\xi})$; this is the route taken in our numerical experiments.

2. The Shrinkage Path Heuristic

We begin by simplifying the feasible set of Problem (1). For any $\rho > 0$ and each point $\xi \in \Xi$, the Wasserstein ball $\mathbb{B}_\rho(\hat{\mathbb{P}})$ contains a distribution that places positive mass at ξ , so the almost-sure constraints decouple from the radius and reduce to the support-only robust requirement that $g_m(\mathbf{x}, \xi) \leq \gamma_m$ for all $\xi \in \Xi$ and every $m \in [M]$. Accordingly, the feasible set of Problem (1) reduces to

$$\mathcal{X}(\Xi) := \{\mathbf{x} \in \mathcal{X} : g_m(\mathbf{x}, \xi) \leq \gamma_m \ \forall \xi \in \Xi, \forall m \in [M]\}. \quad (2)$$

We assume throughout that $\mathcal{X}(\Xi)$ is non-empty, otherwise, Problem (1) is trivially infeasible. Convexity of $\mathcal{X}$ together with convexity of each $g_m(\cdot, \xi)$ ensures that $\mathcal{X}(\Xi)$ is itself convex. The set $\mathcal{X}(\Xi)$ is described by semi-infinite constraints. When each $g_m(\mathbf{x}, \cdot)$ is concave on Ξ and Ξ admits a conic representation, classical robust counterpart techniques yield a tractable finite-dimensional reformulation (Ben-Tal et al. 2009). More broadly, $\mathcal{X}(\Xi)$ remains tractable whenever it admits a polynomial-time separation oracle, that is, whenever for any candidate $\mathbf{x}^*$ a violated constraint can be identified by maximizing $g_m(\mathbf{x}^*, \xi)$ over $\xi \in \Xi$ in polynomial time.

The *sample average approximation* (SAA) solution is any optimal solution $\mathbf{x}(0)$ to

$$\underset{\mathbf{x} \in \mathcal{X}(\Xi)}{\text{minimize}} \quad \mathbb{E}_{\hat{\mathbb{P}}}[\ell(\mathbf{x}, \tilde{\xi})]. \quad (3a)$$

We assume that Problem (3a) attains a finite optimal value, ensuring that $\mathbf{x}(0)$ is well-defined. Since $\hat{\mathbb{P}} \in \mathbb{B}_\rho(\hat{\mathbb{P}})$, this further implies that the WDRO problem (1) is bounded below. Whenever ℓ is convex in $\mathbf{x}$ and $\mathcal{X}(\Xi)$ admits a polynomial-time separation oracle, Problem (3a) is solvable in polynomial time; the Wasserstein ambiguity set plays no role at this extreme.

The *robust optimization* (RO) solution is any optimal solution $\mathbf{x}(\infty)$ to

$$\underset{\mathbf{x} \in \mathcal{X}(\Xi)}{\text{minimize}} \quad \sup_{\xi \in \Xi} \ell(\mathbf{x}, \xi). \quad (3b)$$

We assume that Problem (3b) attains a finite optimal value, ensuring that $\mathbf{x}(\infty)$ is well-defined. This holds, for instance, when Ξ is compact and ℓ is upper semicontinuous in ξ , or more generally when $\ell(\mathbf{x}, \cdot)$ is bounded above on Ξ for every $\mathbf{x} \in \mathcal{X}(\Xi)$. Problem (3b) is tractable when $\ell(\mathbf{x}, \cdot)$ is concave on Ξ . In the convex-in- ξ regime of interest in this paper, the problem is NP-hard in general (Pardalos and Vavasis 1991)—the same source of intractability that afflicts the inner suprema in the dual reformulation of Problem (1). Even so, as in classical robust optimization, this worst-case separation is routinely handled in practice by cutting-plane methods that iteratively add violated worst-case scenarios (Mutapcic and Boyd 2009).

Define the *shrinkage path* as

$$\hat{\mathcal{X}}(\Xi) := \{(1 - \lambda)\mathbf{x}(0) + \lambda\mathbf{x}(\infty) : \lambda \in [0, 1]\}. \quad (4)$$

Because the SAA and RO problems in (3) need not admit unique optimizers, we fix one solution of each; these endpoints $\mathbf{x}(0)$ and $\mathbf{x}(\infty)$ in turn fix the shrinkage path, and we refer henceforth to *the* SAA solution, *the* RO solution, and *the* shrinkage path. The shrinkage path is a tractable proxy for the WDRO solution path $\{\mathbf{x}^*(\rho) : \rho \geq 0\}$, where $\mathbf{x}^*(\rho)$ denotes any optimal solution of Problem (1) at radius ρ . As shown by Hao and Zhang (2025) for robust linear optimization, $\mathbf{x}^*(\rho)$ traces a one-dimensional Bregman-type regularization path; in our setting, however, it is generally intractable to compute. We refer to λ as the *shrinkage parameter*: $\lambda = 0$ corresponds to the SAA solution $\mathbf{x}(0)$, $\lambda = 1$ recovers the RO solution $\mathbf{x}(\infty)$, and intermediate values balance nominal fit against robustness. Since both endpoints lie on the shrinkage path, optimizing over it weakly dominates the *SAA–RO heuristic* of selecting the better of $\mathbf{x}(0)$ and $\mathbf{x}(\infty)$. Figure 1 illustrates the shrinkage path alongside the true WDRO optimizer path and the resulting suboptimality gap.

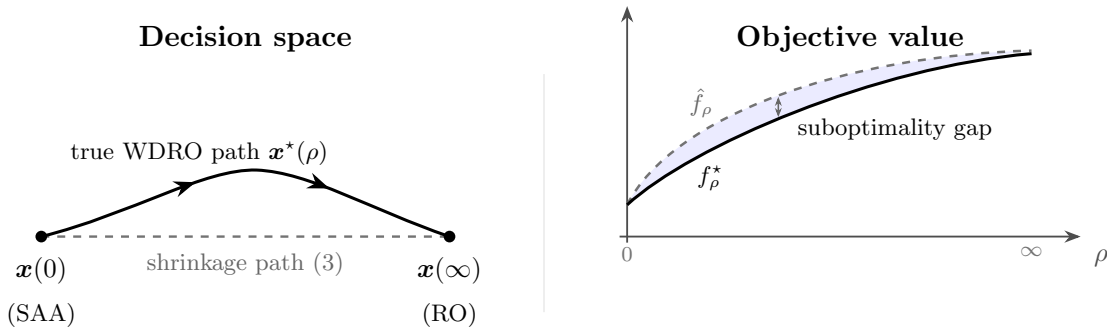


Figure 1 Schematic view of the shrinkage path heuristic. Left: geometric approximation in decision space. The dashed segment is the shrinkage path connecting the SAA and RO endpoints, and the solid curve is the WDRO optimizer path $\mathbf{x}^*(\rho)$. Right: induced value gap as a function of the Wasserstein radius ρ . The dashed curve depicts the shrinkage path values $\hat{f}_\rho$, the solid curve denotes the WDRO optima f_ρ^* , and the shaded region represents the induced suboptimality gaps $\hat{f}_\rho - f_\rho^*$.

Algorithm 1 implements a one-dimensional search over the shrinkage path. Since $\mathcal{X}(\Xi)$ is convex and both endpoints $\mathbf{x}(0)$ and $\mathbf{x}(\infty)$ reside in $\mathcal{X}(\Xi)$, every interpolant $\mathbf{x}$ is feasible in Problem (1).

For a *fixed* radius ρ , Algorithm 1 replaces the joint optimization in Problem (1) with two endpoint problems plus $|\Lambda|$ calls to the worst-case oracle $\text{WC-EVAL}(\cdot; \rho)$. Each oracle call fixes the decision $\mathbf{x}$ and admits the standard Wasserstein dual reformulation (Mohajerin Esfahani and Kuhn 2018, Theorem 4.2), which decomposes the worst-case expectation into N independent k -dimensional subproblems—one per data point, each maximizing the loss minus a penalized transport cost over Ξ . These subproblems are decoupled, parallelize easily, and can be handled by standard global solvers (*e.g.*, branch-and-bound or multistart local search) for moderate k . The in-sample speedup over the joint problem is typically modest, and is not the main benefit of the heuristic. The larger savings

Algorithm 1 Shrinkage Path Heuristic for Problem (1)

Require: Empirical law $\hat{\mathbb{P}} := \frac{1}{N} \sum_{i=1}^N \delta_{\xi_i}$ with samples $\{\hat{\xi}_i\}_{i=1}^N$; radius $\rho \geq 0$; search grid $\Lambda \subset [0, 1]$;
 worst-case evaluation oracle $\text{WC-EVAL}(\mathbf{x}; \rho) = \sup \{ \mathbb{E}_{\mathbb{Q}}[\ell(\mathbf{x}, \tilde{\xi})] : \mathbb{Q} \in \mathbb{B}_{\rho}(\hat{\mathbb{P}}) \}$.

- 1: Initialize incumbent objective $J^* \leftarrow +\infty$ and incumbent solution $\mathbf{x}^* \leftarrow \emptyset$.
- 2: **for** $\lambda \in \Lambda$ **do**
- 3: Compute $\mathbf{x} \leftarrow (1 - \lambda)\mathbf{x}(0) + \lambda\mathbf{x}(\infty)$ and $J \leftarrow \text{WC-EVAL}(\mathbf{x}; \rho)$.
- 4: **if** $J < J^*$ **then**
- 5: $J^* \leftarrow J$ and $\mathbf{x}^* \leftarrow \mathbf{x}$.
- 6: **end if**
- 7: **end for**
- 8: **return** $\mathbf{x}^*$ with objective value J^* .

arise *out-of-sample*, when ρ is itself tuned by K -fold cross-validation. In practice, ρ is rarely chosen a priori; it is selected from a candidate grid $\mathcal{R} \subset \mathbb{R}_+$ to optimize cross-validated performance. For the exact WDRO problem, this requires $K|\mathcal{R}|$ solves of Problem (1), each of which is precisely the joint optimization that motivated the heuristic in the first place. The shrinkage path, by contrast, does not depend on ρ , and only one of its two endpoints depends on the data: the typically harder RO endpoint $\mathbf{x}(\infty)$ is computed once from $\mathcal{X}$, ℓ , and Ξ alone, while the typically easier SAA endpoint $\mathbf{x}(0)$ is recomputed per fold from the corresponding training sample. The entire path is then scored on the validation sample *without any worst-case evaluations*. Concretely, denote by $\hat{\mathbb{P}}^{\text{tr}}$ and $\hat{\mathbb{P}}^{\text{val}}$ the empirical distributions on a fold's training and validation samples, write $\mathbf{x}^{\text{tr}}(0)$ for the SAA endpoint computed from $\hat{\mathbb{P}}^{\text{tr}}$, replace the in-sample oracle in Algorithm 1 with the validation-fold sample average

$$\text{VAL-EVAL}(\mathbf{x}) := \mathbb{E}_{\hat{\mathbb{P}}^{\text{val}}}[\ell(\mathbf{x}, \tilde{\xi})] = \frac{1}{N^{\text{val}}} \sum_{i=1}^{N^{\text{val}}} \ell(\mathbf{x}, \hat{\xi}_i^{\text{val}}),$$

and average VAL-EVAL across the K folds for each λ on the search grid. Feasibility need not be checked separately, since $\mathcal{X}(\Xi)$ does not depend on ρ and every interpolant $(1 - \lambda)\mathbf{x}^{\text{tr}}(0) + \lambda\mathbf{x}(\infty)$ remains in $\mathcal{X}(\Xi)$ by convexity. Let $\lambda^* \in \Lambda$ minimize the cross-validated VAL-EVAL along the path; the final decision $(1 - \lambda^*)\mathbf{x}(0) + \lambda^*\mathbf{x}(\infty)$ is then formed by recomputing the SAA endpoint $\mathbf{x}(0)$ on the full sample $\hat{\mathbb{P}}$, while the RO endpoint $\mathbf{x}(\infty)$ is reused. The heuristic thus replaces cross-validated tuning of ρ , which requires $K|\mathcal{R}|$ solves of Problem (1), with cross-validated tuning of λ , which requires one RO solve, K SAA solves, and sample-average evaluations along the path. The resulting savings are typically several orders of magnitude.

3. A Priori Performance Guarantees

To obtain a priori guarantees, this section focuses on the following structured setting.

ASSUMPTION 1 (Lipschitz Loss). *Problem (1) satisfies the following conditions.*

- (i) *The ground metric is the Euclidean norm: $d(\xi, \xi') = \|\xi - \xi'\|_2$.*
- (ii) *The loss is $\ell(\mathbf{x}, \xi) = L(\xi^\top \mathbf{x})$ with $L: \mathbb{R} \rightarrow \mathbb{R}_+$ convex and Lipschitz with constant $\text{Lip}(L)$.*

Pairing the Euclidean ground metric in (i) with the bilinear coupling $\xi^\top \mathbf{x}$ in (ii) ensures that, for every fixed decision $\mathbf{x}$, the loss $\xi \mapsto L(\xi^\top \mathbf{x})$ is Lipschitz in ξ with the constant $\text{Lip}(L)\|\mathbf{x}\|_2$, which in turn yields explicit, decision-dependent bounds on how fast the worst-case expectation can grow with ρ relative to its nominal counterpart. This implies in particular that the number of uncertain parameters k coincides with the number of decisions n in this section. The Lipschitz assumption on L in (ii) further keeps the worst-case expectation finite when Ξ is unbounded, and convexity of L ensures that the WDRO problem is convex in $\mathbf{x}$. Under these conditions, Problem (1) admits closed-form suboptimality bounds and structural insights that are difficult to obtain in the general case. In particular, when $\Xi = \mathbb{R}^k$, the WDRO problem reduces to a regularized SAA problem (Shafieezadeh-Abadeh et al. 2019, Theorem 4(ii)). In such settings, the heuristic is not needed in practice, since the WDRO problem is itself tractable. We study them because they can be analyzed in closed form, and the resulting guarantees help interpret the strong performance of the heuristic on the harder instances of Section 5.

We define the *additive suboptimality* of a candidate solution at a given Wasserstein radius ρ as the gap between its worst-case expected loss and the optimal value of Problem (1) at that radius. The worst-case additive suboptimality of a heuristic is the supremum of this gap over all $\rho \geq 0$. We say that a bound on this quantity is *tight* when equality is attained by some problem instance, and *asymptotically tight* when equality is approached in the limit along a sequence of instances parameterized by some structural quantity. When neither holds, we say the bound has *slack*.

We first bound the additive suboptimality of our shrinkage path heuristic when Ξ is bounded.

THEOREM 1 (Bounded Support). *Let $\text{diam}(\Xi) = \sup\{\|\xi - \xi'\|_2 : \xi, \xi' \in \Xi\} < \infty$. Then, for any $\rho \in [0, +\infty)$, the additive suboptimality of the shrinkage path heuristic is at most*

$$\rho \left[1 - \frac{\rho}{\text{diam}(\Xi)} \right]_+ \text{Lip}(L)\|\mathbf{x}(0)\|_2,$$

which is maximized at $\rho = \text{diam}(\Xi)/2$. Moreover, the bound is pointwise tight: for every $\rho \in [0, \infty)$, the bound is attained by some problem instance.

The bound in Theorem 1 vanishes at both endpoints of $\rho \in [0, \text{diam}(\Xi)]$ and is largest near the midpoint. Thus, the heuristic is provably near-optimal for small ρ and again as ρ approaches $\text{diam}(\Xi)$. The small- ρ regime is the practically relevant one: the out-of-sample performance of DRO typically improves over SAA at small radii and worsens past a critical threshold (Anderson and

Philpott 2022). The heuristic therefore suits the common goal of improving SAA with a modest amount of distributional robustness.

We now prove an analogous result for the unbounded support case.

PROPOSITION 1 (Unbounded Support). *Let $\Xi = \mathbb{R}^k$, and let $\bar{r} := \frac{1}{N} \sum_{i=1}^N \|\hat{\xi}_i\|_2$ denote the effective diameter of the support. The additive suboptimality of the shrinkage path heuristic is at most*

$$\text{Lip}(L) \left(\rho \left(1 - \frac{\rho}{\bar{r}} \right) \|\mathbf{x}(0)\|_2 + \frac{\rho^2}{\bar{r}} \|\mathbf{x}(\infty)\|_2 \right) \quad \forall \rho \in [0, \bar{r}],$$

which is maximized at $\rho = \bar{r}/2$ when $\mathbf{x}(\infty) = \mathbf{0}$.

Proposition 1 parallels Theorem 1, with the (finite) effective diameter $\bar{r}$ serving as a data-driven surrogate for the (infinite) support diameter $\text{diam}(\Xi)$: below $\bar{r}$, the bound is a weighted sum of two terms—one proportional to the SAA endpoint $\|\mathbf{x}(0)\|_2$ with weight $\rho(1 - \rho/\bar{r})$, and one governed by the RO endpoint $\|\mathbf{x}(\infty)\|_2$ with weight $\rho^2/\bar{r}$ —and it reduces to a parabola with peak at $\rho = \bar{r}/2$ in the case where $\mathbf{x}(\infty) = \mathbf{0}$. Unlike Theorem 1, the bound is not in general tight. The bound can be extended to $\rho > \bar{r}$ and admits a closed-form worst case over all $\rho \geq 0$; we relegate these refinements to the GitHub companion.

The bounds in Theorem 1 and Proposition 1 depend only on coarse global parameters—the Lipschitz constant $\text{Lip}(L)$, the support diameter $\text{diam}(\Xi)$ or its effective counterpart $\bar{r}$, and the endpoint norms $\|\mathbf{x}(0)\|_2$ and $\|\mathbf{x}(\infty)\|_2$. In particular, they do not reflect the curvature of the loss function or the conditioning of the data. We now specialize to a quadratic loss setting that admits sharper, data-dependent guarantees. In addition to an improved additive bound, this setting yields a *multiplicative suboptimality* bound: a bound on the ratio of the worst-case additive suboptimality of the shrinkage path heuristic to that of the SAA–RO heuristic, which selects the better of the two endpoints $\mathbf{x}(0)$ and $\mathbf{x}(\infty)$. This bound shows that the shrinkage path heuristic strictly improves upon the SAA–RO heuristic whenever the data are not perfectly conditioned. This justifies exploring the full shrinkage path rather than restricting attention to its two endpoints.

ASSUMPTION 2 (Quadratic Loss). *Let $\mathcal{Z} := \{\mathbf{x}^\top \hat{\xi}_i : \|\mathbf{x}\|_2 \leq \|\mathbf{x}(0)\|_2, i \in [N]\}$ collect the scalar arguments along the shrinkage path. In addition to Assumption 1, the following conditions hold.*

- (i) *The support and feasible set are unconstrained: $\Xi = \mathbb{R}^k$ and $\mathcal{X}(\Xi) = \mathbb{R}^n$.*
- (ii) *The loss $L : \mathbb{R} \rightarrow \mathbb{R}_+$ is differentiable, and is quadratic on $\mathcal{Z}$, that is, $L(z) = z^2 + bz + c$ on $\mathcal{Z}$ for some $b, c \in \mathbb{R}$.*
- (iii) *The empirical second-moment matrix $\mathbf{H} := \frac{1}{N} \sum_{i=1}^N \hat{\xi}_i \hat{\xi}_i^\top$ is positive definite, with eigenvalues $0 < \lambda_{\min} \leq \dots \leq \lambda_{\max}$ and condition number $\kappa := \lambda_{\max}/\lambda_{\min}$.*

Condition (i) reduces the problem geometry to a form where closed-form analysis is possible: $\Xi = \mathbb{R}^k$ collapses the WDRO problem to a regularized SAA problem with a Euclidean norm penalty, and $\mathcal{X}(\Xi) = \mathbb{R}^n$ forces $\mathbf{x}(\infty) = \mathbf{0}$, so the shrinkage path reduces to the segment $\{\lambda \mathbf{x}(0) : \lambda \in [0, 1]\}$. Condition (ii) is consistent with Assumption 1(ii), provided L extends to a globally Lipschitz function on $\mathbb{R}$. This is satisfied, for example, by the Huber loss $L(z) = z^2$ for $|z| \leq \delta$ and $L(z) = 2\delta|z| - \delta^2$ otherwise: as long as $\mathcal{Z} \subseteq [-\delta, \delta]$, the loss is quadratic on $\mathcal{Z}$ and globally Lipschitz with constant 2δ . Our results readily extend to losses with a continuous, bounded second derivative on $\mathcal{Z}$, with $\lambda_{\min}$ and $\lambda_{\max}$ in the bounds replaced by the local strong convexity and smoothness constants of L on $\mathcal{Z}$; we omit the details for brevity.

The following theorem refines the additive bound of Proposition 1 under the additional structure of Assumption 2.

THEOREM 2 (Additive Suboptimality; Quadratic Loss). *Let Assumption 2 hold. The worst-case additive suboptimality of the shrinkage path heuristic is bounded above by*

$$\lambda_{\max} \|\mathbf{x}(0)\|_2^2 \frac{(\kappa - 1)^2}{\kappa [\varphi^5(\kappa - 1) + 27]}, \quad (5)$$

where $\varphi = (\sqrt{5} + 1)/2$ is the golden ratio. The bound is asymptotically tight as $\kappa \rightarrow 1^+$ and $\kappa \rightarrow \infty$.

The bound (5) factors into two terms: $\lambda_{\max} \|\mathbf{x}(0)\|_2^2$ and a κ -dependent factor. The latter behaves like $(\kappa - 1)^2/27$ to leading order as $\kappa \rightarrow 1^+$, so the shrinkage path is provably near-optimal whenever the empirical second-moment matrix is well-conditioned; it increases monotonically with κ and approaches $1/\varphi^5 \approx 0.0902$ in the severely ill-conditioned limit $\kappa \rightarrow \infty$. The appearance of the golden ratio reflects the limiting worst-case geometry, described in the following corollary.

COROLLARY 1 (Worst-case Geometry). *Let Assumption 2 hold, and let ρ_{gap} , ρ_0^h , and ρ_0^* denote, respectively, a radius at which the suboptimality gap is maximized, the radius at which the shrinkage path solution reaches $\mathbf{0}$, and the smallest radius at which the WDRO solution reaches $\mathbf{0}$. As $\kappa \rightarrow \infty$, the worst-case bound in Theorem 2 is approached by instances for which*

$$\rho_{\text{gap}} : \rho_0^h : \rho_0^* \rightarrow 1 : \varphi : \varphi^2.$$

The corollary describes the configuration that produces the worst case of Theorem 2, in which the two forms of shrinkage are hardest to reconcile. The WDRO problem shrinks eigendirections at non-uniform rates—in the limiting worst case, the $\lambda_{\min}$ direction collapses to zero while the $\lambda_{\max}$ direction shrinks by a factor of $1/\varphi$ —whereas the shrinkage path heuristic applies a single uniform factor across all coordinates. The heuristic incurs its largest gap at the configuration where this mismatch is most costly; there, the uniform factor $1/\varphi^2$ and the dominant per-coordinate factor

$1/\varphi$ sum to one: $\frac{1}{\varphi} + \frac{1}{\varphi^2} = 1$. Thus, no single direction is responsible for the bound: it arises because one common shrinkage factor cannot match the WDRO solution in directions of very different magnitude.

We now turn to the multiplicative suboptimality of our heuristic.

THEOREM 3 (Multiplicative Suboptimality). *Let Assumption 2 hold. The multiplicative suboptimality of the shrinkage path heuristic relative to the SAA-RO heuristic is bounded above by*

$$\frac{(\kappa - 1)^2(\kappa + 1)}{(\kappa^{3/2} + 1)^2}, \quad (6)$$

and the bound is strictly less than one whenever $\kappa \geq 1$.

The bound (6) equals zero at $\kappa = 1$ and increases monotonically to one as $\kappa \rightarrow \infty$, so the shrinkage path offers the largest relative improvement over the SAA-RO heuristic for well- to moderately conditioned data. The bound is generally not tight: it relaxes the worst-case suboptimality of the SAA-RO heuristic to a lower bound, and a smaller denominator enlarges the ratio.

REMARK 2 (CONNECTION TO JAMES-STEIN SHRINKAGE). Under Assumption 2, the solution of the shrinkage path heuristic admits the closed form

$$\hat{\mathbf{x}}(\rho) = \left[1 - \frac{\rho \|\mathbf{x}(0)\|_2}{2 \mathbf{x}(0)^\top \mathbf{H} \mathbf{x}(0)} \right]_+ \mathbf{x}(0),$$

a non-negative scalar multiple of the SAA solution, with the multiplier shrinking from one toward zero as ρ grows. This is formally reminiscent of the positive-part James-Stein estimator $\hat{\boldsymbol{\theta}}^{JS+} = [1 - (n-2)/\|\mathbf{Y}\|_2^2]_+ \mathbf{Y}$ (James and Stein 1961, Baranchik 1970). In the isotropic case $\mathbf{H} = \mathbf{I}$, our shrinkage factor reduces to $1 - \rho/(2\|\mathbf{x}(0)\|_2)$; identifying $\mathbf{Y}$ with $\mathbf{x}(0)$, the heuristic and James-Stein shrinkage factors then coincide at $\rho = 2(n-2)/\|\mathbf{x}(0)\|_2$. The key difference is that the degree of shrinkage is governed by the ambiguity radius ρ rather than the sample size, suggesting that our shrinkage path heuristic implements an ambiguity-driven analogue of classical shrinkage regularization.

4. A Posteriori Performance Guarantees

We next develop computable a posteriori certificates for the suboptimality of our shrinkage path heuristic. Since the shrinkage path is a subset of the WDRO feasible set, the objective value J^* returned by Algorithm 1 *upper* bounds the optimal value of Problem (1). This section complements our upper bound with a matching *lower* bound. Our construction mirrors the primal shrinkage path on the dual side: a minimax interchange first recasts the optimal value of Problem (1) as a maximization over distributions $\mathbb{Q}$ in the Wasserstein ball $\mathbb{B}_\rho(\hat{\mathbb{P}})$, and this maximization is then restricted to the line segment joining the empirical distribution $\hat{\mathbb{P}}$ to a worst-case distribution at $\rho = \infty$. To secure the minimax interchange and the existence of a worst-case distribution at $\rho = \infty$, this section imposes the following regularity assumption on Problem (1).

ASSUMPTION 3 (A Posteriori Regularity). *Problem (1) satisfies the following conditions.*

- (i) *The SAA objective $\mathbf{x} \mapsto \mathbb{E}_{\hat{\mathbb{P}}}[\ell(\mathbf{x}, \tilde{\boldsymbol{\xi}})]$ is coercive on $\mathcal{X}$.*
- (ii) *There exist a neighborhood $\mathcal{U}$ of the RO solution $\mathbf{x}(\infty)$ and a compact set $\mathcal{K} \subseteq \Xi$ such that*

$$\sup_{\boldsymbol{\xi} \in \Xi} \ell(\mathbf{x}, \boldsymbol{\xi}) = \max_{\boldsymbol{\xi} \in \mathcal{K}} \ell(\mathbf{x}, \boldsymbol{\xi}) \quad \forall \mathbf{x} \in \mathcal{U} \cap \mathcal{X}.$$

Recall that a function $f : \mathcal{X} \rightarrow \mathbb{R} \cup \{+\infty\}$ is *coercive* on $\mathcal{X}$ if all its sublevel sets on $\mathcal{X}$ are bounded, or equivalently, if $f(\mathbf{x}_k) \rightarrow +\infty$ along every sequence $\{\mathbf{x}_k\} \subseteq \mathcal{X}$ with $\|\mathbf{x}_k\| \rightarrow \infty$. Condition (i) secures the minimax interchange between the decisions $\mathbf{x}$ and the distributions $\mathbb{Q}$ in Problem (1), while Condition (ii) guarantees the existence of a worst-case distribution at $\rho = \infty$ that is supported on finitely many atoms. The latter condition holds, for example, whenever Ξ is compact.

Denote by $J^*(\rho)$ the optimal value of Problem (1) at radius ρ . We henceforth denote by $\bar{J}(\rho)$ the optimum of Problem (1) restricted to the full shrinkage path $\hat{\mathcal{X}}(\Xi)$, with the shrinkage parameter λ ranging over the full interval $[0, 1]$. In the notation of Section 3, $J^*(\rho) = f_\rho^*$ is the WDRO optimum and $\bar{J}(\rho) = \hat{f}_\rho$ is the value of the shrinkage path heuristic. The objective value returned by Algorithm 1 is a finite-grid approximation of $\bar{J}(\rho)$, since Algorithm 1 searches only over the grid $\Lambda \subset [0, 1]$. The full-path optimum itself is given by

$$\bar{J}(\rho) := \min_{\mathbf{x} \in \hat{\mathcal{X}}(\Xi)} \sup_{\mathbb{Q} \in \mathbb{B}_\rho(\hat{\mathbb{P}})} \mathbb{E}_{\mathbb{Q}}[\ell(\mathbf{x}, \tilde{\boldsymbol{\xi}})] \geq J^*(\rho),$$

where the inequality follows from $\hat{\mathcal{X}}(\Xi) \subseteq \mathcal{X}(\Xi)$. To certify the resulting approximation gap $\bar{J}(\rho) - J^*(\rho)$ without access to $J^*(\rho)$, we construct a matching *lower* bound $\underline{J}(\rho)$ via a dual shrinkage path that mirrors the primal one. The construction is motivated by the minimax interchange

$$J^*(\rho) = \min_{\mathbf{x} \in \mathcal{X}(\Xi)} \sup_{\mathbb{Q} \in \mathbb{B}_\rho(\hat{\mathbb{P}})} \mathbb{E}_{\mathbb{Q}}[\ell(\mathbf{x}, \tilde{\boldsymbol{\xi}})] = \sup_{\mathbb{Q} \in \mathbb{B}_\rho(\hat{\mathbb{P}})} \inf_{\mathbf{x} \in \mathcal{X}(\Xi)} \mathbb{E}_{\mathbb{Q}}[\ell(\mathbf{x}, \tilde{\boldsymbol{\xi}})], \quad (7)$$

which holds for every $\rho \in \mathbb{R}_+ \cup \{+\infty\}$ under Assumption 3(i) (Kuhn et al. 2025, Corollary 5.16). Note that the inner infimum on the right-hand side may not be attained because, for a fixed $\mathbb{Q}$, the map $\mathbf{x} \mapsto \mathbb{E}_{\mathbb{Q}}[\ell(\mathbf{x}, \tilde{\boldsymbol{\xi}})]$ need not be coercive. This min-sup = sup-inf identity instantiates the *primal worst equals dual best* principle (Beck and Ben-Tal 2009, Zhen et al. 2025); the interchange and the existence of an associated worst-case (least-favorable) distribution are well studied in this setting (Blanchet and Murthy 2019).

The dual shrinkage path is the line segment between the two saddle distributions at $\rho = 0$ and $\rho = \infty$. At $\rho = 0$, the Wasserstein ball $\mathbb{B}_\rho(\hat{\mathbb{P}})$ collapses to the empirical distribution $\hat{\mathbb{P}}$. At $\rho = \infty$, we fix any optimizer of the dual maximization in (7) as $\rho \rightarrow \infty$, namely

$$\mathbb{Q}(\infty) \in \arg \max_{\mathbb{Q} \in \mathcal{P}(\Xi)} \inf_{\mathbf{x} \in \mathcal{X}(\Xi)} \mathbb{E}_{\mathbb{Q}}[\ell(\mathbf{x}, \tilde{\boldsymbol{\xi}})], \quad (8)$$

and call it the *RO distribution* $\mathbb{Q}(\infty)$, paralleling the RO solution $\mathbf{x}(\infty)$ on the primal side. The GitHub companion establishes that under Assumption 3, such an optimizer exists and can be chosen to be supported on at most $n + 1$ atoms. Fixing $\mathbb{Q}(\infty)$ in turn fixes the dual shrinkage path; although both are determined only up to the choice of optimizer, we henceforth speak of *the* RO distribution and *the* dual shrinkage path. The latter is defined as

$$\widehat{\mathcal{B}}(\hat{\mathbb{P}}) := \{(1 - \lambda)\hat{\mathbb{P}} + \lambda\mathbb{Q}(\infty) : \lambda \in [0, 1]\}.$$

The primal and dual shrinkage paths are symmetric: the primal path $\widehat{\mathcal{X}}(\Xi)$ is the line segment between the SAA solution $\mathbf{x}(0)$ and the RO solution $\mathbf{x}(\infty)$, while the dual path $\widehat{\mathcal{B}}(\hat{\mathbb{P}})$ is the line segment between the empirical distribution $\hat{\mathbb{P}}$ and the RO distribution $\mathbb{Q}(\infty)$. By the minimax interchange (7), the endpoints of the two paths pair up: $(\mathbf{x}(0), \hat{\mathbb{P}})$ at $\rho = 0$ and $(\mathbf{x}(\infty), \mathbb{Q}(\infty))$ at $\rho = \infty$. Both pairs are in fact Nash equilibria of the game in which the decision maker minimizes and nature maximizes $\mathbb{E}_{\mathbb{Q}}[\ell(\mathbf{x}, \tilde{\boldsymbol{\xi}})]$: at $\rho = 0$ the ambiguity set collapses to the singleton $\{\hat{\mathbb{P}}\}$, leaving neither player a profitable deviation, while at $\rho = \infty$ the minimax interchange (7) equates the min-sup and max-inf values, so their outer optimizers $\mathbf{x}(\infty)$ and $\mathbb{Q}(\infty)$ are mutual best responses.

Restricting the dual maximization in (7) from the Wasserstein ball $\mathbb{B}_{\rho}(\hat{\mathbb{P}})$ to its intersection with the dual shrinkage path $\widehat{\mathcal{B}}(\hat{\mathbb{P}})$ yields the matching a posteriori lower bound

$$\underline{J}(\rho) := \max_{\mathbb{Q} \in \widehat{\mathcal{B}}(\hat{\mathbb{P}}) \cap \mathbb{B}_{\rho}(\hat{\mathbb{P}})} \inf_{\mathbf{x} \in \mathcal{X}(\Xi)} \mathbb{E}_{\mathbb{Q}}[\ell(\mathbf{x}, \tilde{\boldsymbol{\xi}})] \leq J^*(\rho). \quad (9)$$

The a posteriori gap $\bar{J}(\rho) - \underline{J}(\rho)$ then provides a suboptimality certificate for the shrinkage path heuristic in Algorithm 1.

The lower bound $\underline{J}(\rho)$ admits a reduction to a finite-scenario stochastic optimization problem.

THEOREM 4 (Dual Shrinkage Closed Form). *Whenever $\hat{\mathbb{P}} \neq \mathbb{Q}(\infty)$, the a posteriori lower bound (9) admits the finite-scenario reformulation*

$$\underline{J}(\rho) = \min_{\mathbf{x} \in \mathcal{X}(\Xi)} \mathbb{E}_{\hat{\mathbb{Q}}(\rho)}[\ell(\mathbf{x}, \tilde{\boldsymbol{\xi}})] \quad \text{with} \quad \hat{\mathbb{Q}}(\rho) = \mu(\rho)\hat{\mathbb{P}} + (1 - \mu(\rho))\mathbb{Q}(\infty) \quad (10)$$

under the mixture $\mu(\rho) := [1 - \rho/W(\hat{\mathbb{P}}, \mathbb{Q}(\infty))]_{+}$. In the degenerate case where $\hat{\mathbb{P}} = \mathbb{Q}(\infty)$, the dual shrinkage path collapses to $\{\hat{\mathbb{P}}\}$, and (10) holds with $\hat{\mathbb{Q}}(\rho) = \hat{\mathbb{P}}$ for every $\rho \geq 0$.

Theorem 4 identifies $\hat{\mathbb{Q}}(\rho)$ as the distribution on the dual shrinkage path that lies as far from $\hat{\mathbb{P}}$ as the radius permits: it sits on the boundary of the Wasserstein ball, $W(\hat{\mathbb{Q}}(\rho), \hat{\mathbb{P}}) = \rho$, whenever $\rho \leq W(\hat{\mathbb{P}}, \mathbb{Q}(\infty))$, and coincides with $\mathbb{Q}(\infty)$ otherwise; equivalently, $W(\hat{\mathbb{Q}}(\rho), \hat{\mathbb{P}}) = \min\{\rho, W(\hat{\mathbb{P}}, \mathbb{Q}(\infty))\}$. Geometrically, the dual shrinkage path $\widehat{\mathcal{B}}(\hat{\mathbb{P}})$ stays fixed while the Wasserstein ball $\mathbb{B}_{\rho}(\hat{\mathbb{P}})$ expands with ρ , so $\hat{\mathbb{Q}}(\rho)$ slides along the path at the ball's boundary until it reaches $\mathbb{Q}(\infty)$. Like the primal

shrinkage path, the dual path does not vary with ρ : evaluating $\underline{J}(\rho)$ requires only a one-off RO distribution construction, a one-off Wasserstein-distance evaluation between two finitely supported distributions, and a single finite-scenario stochastic optimization per radius, as summarized in Algorithm 2. Sweeping a grid of radii therefore traces the entire lower-bound curve at negligible additional cost per radius. Together with $\bar{J}(\rho)$, this brackets $J^*(\rho)$ from above and below, and so certifies the quality of the heuristic at every radius without solving Problem (1) exactly.

Algorithm 2 A Posteriori Lower Bound Computation

Require: Empirical law $\hat{\mathbb{P}} := \frac{1}{N} \sum_{i=1}^N \delta_{\hat{\xi}_i}$ with samples $\{\hat{\xi}_i\}_{i=1}^N$; radius $\rho \geq 0$.

- 1: Construct the RO distribution $\mathbb{Q}(\infty)$ as defined in (8).
 - 2: Compute $W(\hat{\mathbb{P}}, \mathbb{Q}(\infty))$.
 - 3: Form the dual shrinkage distribution $\hat{\mathbb{Q}}(\rho)$ according to (10).
 - 4: Solve $\underline{J}(\rho) \leftarrow \min_{\mathbf{x} \in \mathcal{X}(\Xi)} \mathbb{E}_{\hat{\mathbb{Q}}(\rho)}[\ell(\mathbf{x}, \tilde{\xi})]$.
 - 5: **return** $\underline{J}(\rho)$.
-

The following theorem establishes further structural properties shared by $\bar{J}(\rho)$, $J^*(\rho)$, and $\underline{J}(\rho)$.

THEOREM 5 (A Posteriori Guarantee). *The a posteriori upper bound $\bar{J}(\rho)$, the WDRO optimum $J^*(\rho)$, and the a posteriori lower bound $\underline{J}(\rho)$ satisfy:*

- (i) **Monotonicity:** *All three values are nondecreasing in ρ .*
- (ii) **Concavity:** *All three values are concave in ρ .*
- (iii) **Endpoint tightness:** *$\bar{J}(\rho) = J^*(\rho) = \underline{J}(\rho)$ at $\rho = 0$ and for every $\rho \geq W(\hat{\mathbb{P}}, \mathbb{Q}(\infty))$.*

Theorem 5 shows that the a posteriori bounds and the exact WDRO value share the same shape: all three start from $J(0)$ (the optimal value of the SAA problem (3a)), rise concavely with ρ , and saturate at $J(\infty)$ (the optimal value of the RO problem (3b)) by $\rho = W(\hat{\mathbb{P}}, \mathbb{Q}(\infty))$.

Algorithm 2 requires constructing the RO distribution $\mathbb{Q}(\infty)$, defined in (8). The following proposition provides three such constructions, with details deferred to Section EC.1.

PROPOSITION 2 (RO Distribution Construction). *Under the standing assumptions, a finitely supported RO distribution $\mathbb{Q}(\infty)$ can be recovered from a finite-dimensional convex problem in each of the following settings.*

- (i) **Convex loss, polyhedral support.** *The support Ξ is a pointed polyhedron, and $\ell(\mathbf{x}, \cdot)$ is convex on Ξ with nonpositive recession along every extreme ray for every $\mathbf{x} \in \mathcal{X}(\Xi)$.*
- (ii) **Finite-max refinement.** *In addition to (i), $\mathcal{X}(\Xi)$ is polyhedral and $\ell = \max_{l \in [L]} \ell_l$, with each $\ell_l(\cdot, \xi)$ convex and differentiable. The resulting distribution has at most $n + 1$ atoms.*

(iii) Convex–concave loss. The support Ξ is compact and $\ell = \max_{l \in [L]} \ell_l$, where every ℓ_l is closed convex with full domain in $\mathbf{x}$ and concave and upper semicontinuous in $\boldsymbol{\xi}$ on Ξ .

Construction (i) yields an RO distribution that is supported on the vertices of Ξ . Construction (ii) specializes (i) and obtains $\mathbb{Q}(\infty)$ from any basic feasible solution of a linear program by exploiting the finite-max differentiable structure of the loss. Construction (iii), finally, recovers $\mathbb{Q}(\infty)$ from a Fenchel dual of the RO problem (3b) and applies to the finite-max convex–concave structure underlying tractable convex reformulations of Wasserstein DRO (Mohajerin Esfahani and Kuhn 2018, Shafiee et al. 2026). The constructions can differ in support size: the GitHub companion proves that an RO distribution on at most $n + 1$ atoms always exists and describes the post-processing routine used to sparsify constructions (i) and (iii) to $n + 1$ atoms. Construction (ii) meets this bound directly.

5. Numerical Experiments

In this section, we evaluate the effectiveness of our shrinkage path heuristic in both a multi-item newsvendor (Section 5.1) and an appointment scheduling setting (Section 5.2). All experiments were conducted on a high-performance computing environment, which hosts AMD EPYC 7742 processors with 64 GB of RAM and runs Python 3.11, CVXPY 1.7.3, MOSEK 11.0.29, and Gurobi 12.0.3. All reported runtimes correspond to single-threaded execution. All solver parameters are set to their default values except where explicitly stated otherwise. We evaluate our shrinkage heuristic along three dimensions: the quality of its a posteriori guarantees, its out-of-sample performance compared to unregularized and regularized SAA (Kleywegt et al. 2002), Wasserstein DRO (Kuhn et al. 2019), and Sinkhorn DRO (Wang et al. 2026), and its scalability relative to these methods. Supplementary implementation details are provided in the GitHub companion.

5.1. Multi-Item Newsvendor

Consider the multi-item newsvendor problem, a classical inventory control model in which the decision maker selects order quantities $\mathbf{x}$ under uncertain demand $\tilde{\boldsymbol{\xi}}$ supported on the box $\Xi = [\mathbf{0}, \mathbf{U}]$ for some $\mathbf{U} \in \mathbb{R}_{++}^n$, and then incurs holding costs for excess inventory and backlogging costs for unmet demand (Zipkin 2000). Given an empirical distribution $\hat{\mathbb{P}} = \frac{1}{N} \sum_{j=1}^N \delta_{\boldsymbol{\xi}_j}$ and a type-1 Wasserstein ball with radius ρ and ℓ_2 ground metric, we benchmark the performance of different methods for (approximately) solving the WDRO newsvendor problem (Mohajerin Esfahani and Kuhn 2018, Gao and Kleywegt 2023)

$$\begin{aligned} & \underset{\mathbf{x}}{\text{minimize}} && \sup_{\mathbb{Q} \in \mathbb{B}_\rho(\hat{\mathbb{P}})} \mathbb{E}_{\mathbb{Q}}[\mathbf{h}^\top [\mathbf{x} - \tilde{\boldsymbol{\xi}}]_+ + \mathbf{b}^\top [\tilde{\boldsymbol{\xi}} - \mathbf{x}]_+] \\ & \text{subject to} && \mathbf{0} \leq \mathbf{x} \leq \mathbf{V}. \end{aligned} \tag{11}$$

where $[\cdot]_+ := \max\{\cdot, 0\}$ is applied componentwise, and $\mathbf{h}, \mathbf{b}, \mathbf{V} \in \mathbb{R}_{++}^n$ denote the marginal holding costs, the marginal backlogging costs, and the order capacities, respectively.

The newsvendor problem is a useful diagnostic because both endpoints of the shrinkage path are available in closed form, while the full WDRO problem remains practically tractable. The *critical ratio* $q_i := b_i / (h_i + b_i) \in (0, 1)$ —the cost-optimal service level (order fractile) of the classical newsvendor—increases in the backlogging cost b_i and decreases in the holding cost h_i , so costlier shortages call for higher order quantities. Let $\widehat{\xi}_{(k),i}$ denote the k th order statistic of the i th demand coordinate, i.e., the k th smallest of its N samples. Then, at $\rho = 0$, the SAA endpoint is the empirical q_i -quantile of demand—the cost-optimal service level of the nominal newsvendor, clipped at capacity, $x_i(0) = \min\{\widehat{\xi}_{(\lceil q_i N \rceil),i}, V_i\}$, $i \in [n]$. As $\rho \rightarrow \infty$, the WDRO problem reduces to the robust problem $\min_{\mathbf{x} \in [0, \mathbf{V}]} \sum_{i=1}^n \max\{h_i x_i, b_i(U_i - x_i)\}$, whose solution is $x_i(\infty) = \min\{q_i U_i, V_i\}$, $i \in [n]$. Finally, for finite ρ , Problem (11) admits an exact conic reformulation (Mohajerin Esfahani and Kuhn 2018, Corollary 5.1): the loss can be written as the maximum of 2^n affine functions of $\boldsymbol{\xi}$, and the box support and ℓ_2 ground metric yield a second-order cone program. For completeness, we provide this SOC reformulation in the GitHub companion.

Throughout this section, all experiments draw from a common master instance of $n = 20$ items. We generate the cost and support parameters by tiling a length-five pattern four times: the holding costs are $h_i = 1$ for all i , the backlogging costs $\mathbf{b}$ and support upper bounds $\mathbf{U}$ cycle through $(2, 6, 8, 10, 10)$ and $(10, 8, 6, 4, 2)$ respectively, and the order capacities are $V_i = 0.6 U_i$; each pattern thus repeats across the four consecutive blocks of five items. Demands are sampled independently as $\tilde{\xi}_i = U_i \tilde{Z}_i$ with $\tilde{Z}_i \sim \text{Beta}(\alpha_i, \beta_i)$, where (α_i, β_i) cycles through $(0.2, 2), (0.2, 1), (0.2, 3), (0.4, 2), (0.3, 2)$, and the empirical distribution $\hat{\mathbb{P}}$ is built from $N = 15$ i.i.d. samples. Experiments that use fewer than 20 items take the data of the first n items.

We first benchmark our a posteriori guarantees with $n = 15$ items, evaluating the worst-case expected loss along a logarithmic grid of 30 values of ρ between 10^{-2} and 10^2 . Figure 2 (left) plots the a posteriori envelopes of Section 4 and the exact WDRO objective and demonstrates that our shrinkage heuristic is near-optimal across all radii and exact at the two endpoints. In addition, the right panel of Figure 2 reports the average per- ρ runtime of the primal approximation, the dual approximation, and the exact WDRO solve.

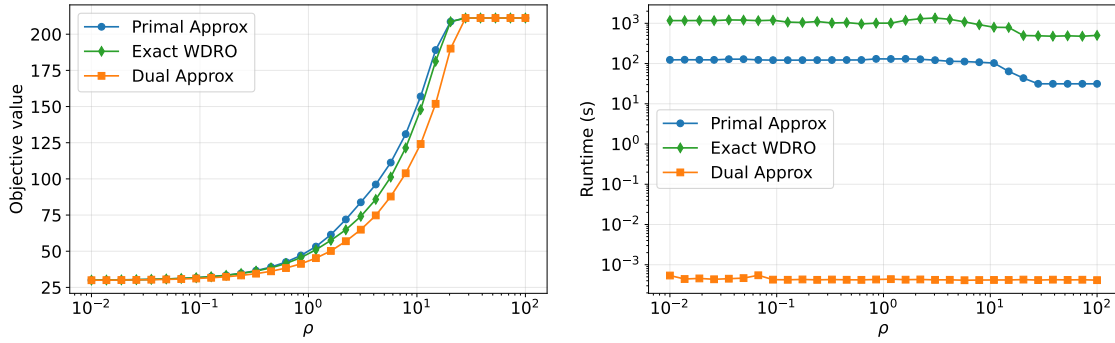


Figure 2 A posteriori upper bound (primal approximation), exact WDRO objective, and a posteriori lower bound (dual approximation) (left panel), and runtimes (right panel) for the 15-item newsvendor problem.

Next, we evaluate the out-of-sample performance of the following methods. First, we include SAA. Second, we include two regularized SAA benchmarks, “ ℓ_1 -RSAA-CV” and “ ℓ_2 -RSAA-CV”, which solve $\min_{\mathbf{x} \in \mathcal{X}} \frac{1}{N} \sum_{j=1}^N \ell(\mathbf{x}, \hat{\xi}_j) + \nu \|\mathbf{x}\|_p^p$ for $p = 1$ and $p = 2$, respectively, where $\nu \geq 0$ is a regularization weight tuned by five-fold cross-validation over 25 values logarithmically spaced between 10^{-3} and 10^2 and augmented with $\nu = 0$. In view of the well-documented connections between Wasserstein DRO and norm regularization (Shafieezadeh-Abadeh et al. 2019, Gao et al. 2024), these benchmarks test whether shrinking the decision per se already confers the benefits of distributional robustness, or whether the direction of shrinkage toward the RO solution matters. Third, we solve the exact WDRO reformulation introduced above, with ρ tuned by five-fold cross-validation over 30 values logarithmically spaced between 0.1 and 15 and augmented with $\rho = 0$; we refer to this approach as “WDRO-CV”. We additionally solve the exact WDRO problem with the cutting-plane algorithm described in the GitHub companion, denoted “WDRO-CV (CP)”; since it returns the same solutions as WDRO-CV, we report its out-of-sample performance only for instances that WDRO-CV cannot solve within our runtime and memory budget, but report its runtime for all n . Fourth, we implement the Sinkhorn DRO method of Wang et al. (2026), which we refer to as “SDRO-CV”, with the nominal radius $\bar{\rho}$ and entropic strength ϵ jointly tuned by five-fold cross-validation over the $22 \times 14 = 308$ candidate $(\bar{\rho}, \epsilon)$ pairs of their default newsvendor configuration (implementation details are provided in the GitHub companion). Fifth, we evaluate our shrinkage path heuristic, with λ tuned by five-fold cross-validation over 30 uniformly spaced values in $[0, 1]$; we refer to this approach as “Shrinkage-CV”. Sixth, “SAA-RO-CV” uses five-fold cross-validation to choose the better of the SAA and RO decisions. Finally, we compute oracle counterparts of the shrinkage heuristic and of exact WDRO, with λ and ρ tuned directly on the test samples; these oracles, “Shrinkage-Oracle” and “WDRO-Oracle”, give the best out-of-sample performance each approach can attain.

When the training and test data are drawn from the same distribution, SAA is already a strong benchmark for the data-driven newsvendor problem. The value of distributional robustness surfaces primarily under a distribution shift. Accordingly, to assess robustness to such a shift between training and deployment, the test data follow the coordinatewise contaminated mixture of Hanasusanto et al. (2015),

$$\mathbb{P}_{\text{test}} = \bigotimes_{i=1}^n \left[(1-c)\mathbb{P}_{\text{train},i} + c\text{Unif}[0, U_i] \right],$$

where $\mathbb{P}_{\text{train},i}$ is the i th training marginal, $\text{Unif}[0, U_i]$ is the uniform distribution on $[0, U_i]$, and the contamination level is $c = 0.2$. We measure out-of-sample performance through the suboptimality of each method, that is, the percentage by which the expected cost of its decision on 50,000 test samples exceeds the smallest expected cost attainable by any feasible decision under the same samples.

Figure 3 reports the out-of-sample suboptimality and average runtime across 100 random seeds, as the number of items varies over $n \in \{8, 10, 12, 15, 20\}$ with the sample size fixed at $N = 15$. The left panel reports out-of-sample suboptimality. The boxes span interquartile ranges, the whiskers reach the most extreme seeds within 1.5 interquartile ranges of the boxes, and seeds beyond the whiskers are omitted. SAA and the two regularized SAA benchmarks perform worst across all problem sizes. Both ℓ_1 -RSAA-CV and ℓ_2 -RSAA-CV underperform SAA; since $\nu = 0$ belongs to the tuning grid, we attribute this shortfall to cross-validation error under the small sample size $N = 15$. Wasserstein DRO and Sinkhorn DRO are marginally better than the shrinkage path heuristic for $n \leq 15$. At $n = 20$, the shrinkage path heuristic outperforms the cutting-plane variant of Wasserstein DRO but remains marginally worse than Sinkhorn DRO. The SAA-RO endpoint heuristic (SAA-RO-CV) also performs among the worst methods across all problem sizes. Overall, for the multi-item newsvendor problem, the shrinkage path heuristic captures around 85%–110% of the median out-of-sample benefit of Wasserstein DRO over SAA. The percentages are computed from the median relative suboptimalities reported in Figure 3: at each n , the gap between SAA and the shrinkage path heuristic is divided by the gap between SAA and Wasserstein DRO, and the range is taken over n .

The right panel of Figure 3 reports the runtimes for each method. Here, the (barely visible) shaded bands cover one standard deviation on either side of the mean runtime. Results for Wasserstein DRO and Sinkhorn DRO are reported only when the corresponding solve finishes within the 72-hour runtime budget. The exact reformulation of Wasserstein DRO exhausts our 64 GB memory budget at $n = 20$, where only its cutting-plane variant remains tractable. The shrinkage path heuristic runs three orders of magnitude faster than either Wasserstein DRO variant or Sinkhorn DRO. These computational savings allow the heuristic to achieve its competitive out-of-sample performance at a fraction of the computational cost and make it especially attractive in large-scale settings.

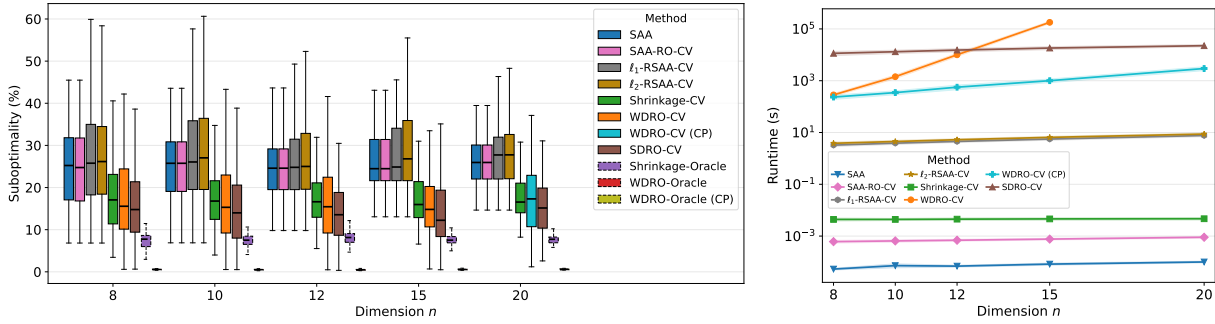


Figure 3 Out-of-sample suboptimality (left) and average runtime (right) of the different methods on the multi-item newsvendor problem with $N = 15$ samples, as the number of items varies over $n \in \{8, 10, 12, 15, 20\}$. The shrinkage path heuristic captures most of the out-of-sample benefit of Wasserstein or Sinkhorn DRO over SAA, at orders of magnitude less computational cost.

5.2. Appointment Scheduling

We consider an appointment scheduling problem in the spirit of Denton and Gupta (2003) and Jiang et al. (2019). A service provider schedules n consecutive appointment slots within a session of maximum duration D and must fix the slot lengths $\mathbf{x} \in \mathbb{R}_+^n$ before observing the uncertain service durations $\tilde{\xi}$. The durations are supported on the box $\Xi = [\underline{\xi}, \bar{\xi}]$. Given the empirical distribution $\hat{\mathbb{P}} = \frac{1}{N} \sum_{j=1}^N \delta_{\xi_j}$, we study the type-1 Wasserstein DRO problem with ℓ_2 ground metric

$$\begin{aligned} & \underset{\mathbf{x}}{\text{minimize}} && \sup_{\mathbb{Q} \in \mathbb{B}_\rho(\hat{\mathbb{P}})} \mathbb{E}_{\mathbb{Q}}[\ell(\mathbf{x}, \tilde{\xi})] \\ & \text{subject to} && \mathbf{x} \geq \mathbf{0}, \quad \sum_{t \in [n]} x_t \leq D, \end{aligned} \quad (12a)$$

where $\ell(\mathbf{x}, \xi)$ is the optimal value of the second-stage problem

$$\begin{aligned} & \underset{\mathbf{w}, \mathbf{s}, o}{\text{minimize}} && \sum_{t=2}^n (c_w w_t^2 + c_s s_t^2) + c_o o^2 \\ & \text{subject to} && w_1 = s_1 = 0, \\ & && w_{t+1} \geq w_t + \xi_t - x_t, \quad w_{t+1} \geq 0 \quad \forall t \in [n-1], \\ & && s_{t+1} \geq x_t - w_t - \xi_t, \quad s_{t+1} \geq 0 \quad \forall t \in [n-1], \\ & && o \geq w_n + \xi_n - x_n, \quad o \geq 0. \end{aligned} \quad (12b)$$

Here, the recourse variables w_t and s_t denote, respectively, the waiting time of patient t and the server's idle time immediately before appointment t , while o is the overtime beyond the scheduled session end $\sum_{t \in [n]} x_t$. The constants $c_w, c_s, c_o > 0$ are the penalty weights on the squared waiting times, idle times, and overtime. Since the second-stage objective is jointly convex in $(\mathbf{x}, \xi, \mathbf{w}, \mathbf{s}, o)$ and the constraints are jointly affine, $\ell(\mathbf{x}, \xi)$ is jointly convex in $(\mathbf{x}, \xi)$.

Unlike the newsvendor loss, which is a maximum of finitely many affine functions of ξ , the appointment scheduling loss $\ell(\mathbf{x}, \xi)$ is a maximum of *infinitely many* affine functions. The extreme-point reduction that would render it finite fails because the second-stage dual is concave-quadratic, rather than linear, in its multipliers (cf. the GitHub companion). Since no finite, exact, convex reformulation is available, we solve the exact WDRO problem using the same cutting-plane algorithm as in the previous subsection, which we denote “WDRO-CP”. As the worst-case evaluation oracle, WC-EVAL($\cdot; \rho$), (cf. Section 2) does not admit an exact finite reformulation either, we approximate each oracle call, including the separation problems within WDRO-CP, by a multi-start alternating maximization procedure (see the GitHub companion for more details). The shrinkage path endpoints, by contrast, remain practically tractable. Writing $\mathcal{X} = \{\mathbf{x} \in \mathbb{R}_+^n : \sum_{t=1}^n x_t \leq D\}$, the SAA and RO endpoints solve the convex quadratically constrained programs

$$\begin{aligned} \mathbf{x}(0) &\in \arg \min_{\mathbf{x} \in \mathcal{X}, \tau \in \mathbb{R}^N} \left\{ \frac{1}{N} \sum_{j=1}^N \tau_j : \ell(\mathbf{x}, \hat{\xi}_j) \leq \tau_j \quad \forall j \in [N] \right\}, \\ \mathbf{x}(\infty) &\in \arg \min_{\mathbf{x} \in \mathcal{X}, \tau \in \mathbb{R}} \left\{ \tau : \ell(\mathbf{x}, \mathbf{v}_k) \leq \tau \quad \forall k \in [2^n] \right\}, \end{aligned}$$

where $\mathbf{v}_k$ ranges over the vertices of Ξ . Optimizing over these vertices suffices because $\xi \mapsto \ell(\mathbf{x}, \xi)$ is convex, so its maximum over the box Ξ is attained at a vertex. Since enumerating all 2^n vertices is impractical, we compute $\mathbf{x}(\infty)$ using a cutting-plane method. In particular, we solve a relaxation over a subset of the vertex constraints and add the most violated one at each iteration, until no vertex constraint is violated.

Throughout this section, all experiments draw from a common master instance of $n = 20$ appointments. We set the support bounds to $\underline{\xi} = \mathbf{0}$ and $\bar{\xi} = (2, \dots, 2)$, the session length to $D = \sum_{i=1}^n (\underline{\xi}_i + \bar{\xi}_i)/2$, and the penalty weights to $c_w = 1.0$ and $c_s = c_o = 0.5$. Service durations are sampled as $\tilde{\xi}_t = \bar{\xi}_t \tilde{Z}_t$, where each $\tilde{Z}_t \in [0, 1]$ follows a two-component Beta mixture. Its shape parameters (α_t^s, β_t^s) cycle through $\{(0.2, 2), (0.2, 1), (0.2, 3), (0.4, 2), (0.3, 2)\}$, with $\alpha_t^s < \beta_t^s$ so that $\text{Beta}(\alpha_t^s, \beta_t^s)$ favors short durations. We draw $\tilde{\mathbf{Z}} \sim (1 - c)\text{Beta}(\boldsymbol{\alpha}^s, \boldsymbol{\beta}^s) + c\text{Beta}(\boldsymbol{\beta}^s, \boldsymbol{\alpha}^s)$ with $c = 0.05$. The parameter-swapped second component favors long durations, so sessions have a 5% chance of running long. The empirical distribution $\hat{\mathbb{P}}$ is built from $N = 10$ i.i.d. samples. Experiments that use fewer than 20 appointments take the data of the first n .

We first construct the a posteriori envelopes of Section 4 for the full $n = 20$ instance and evaluate the worst-case expected loss along a geometric grid of 30 values of ρ between 10^{-3} and 30. Figure 4 reports the resulting primal and dual approximation bounds on the optimal value of (12a) and demonstrates that our shrinkage path heuristic is near-optimal across all radii and exact at the two endpoints. For comparison, it also reports the objective value of the exact WDRO problem, solved

by WDRO-CP, together with the runtime of all three methods, showing that the dual bound is two orders of magnitude cheaper than the primal.

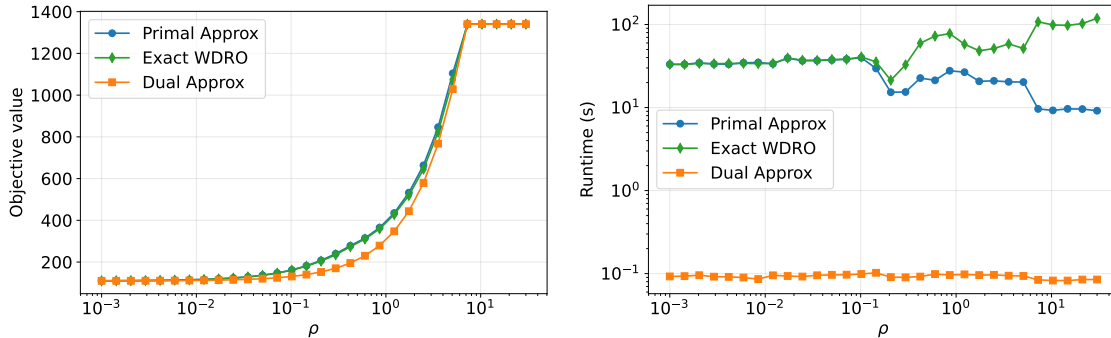


Figure 4 A posteriori upper and lower bounds (left) and average runtimes (right) with varying ρ for the appointment scheduling WDRO problem with $n = 20$ appointments and $N = 10$ samples.

Next, we compare the out-of-sample performance of our shrinkage path heuristic against the same set of methods as in the newsvendor experiment (Section 5.1), with Wasserstein DRO now solved by WDRO-CP. All hyperparameters are again tuned by five-fold cross-validation, and the number of appointments is swept over $n \in \{8, 10, 12, 15, 20\}$. The setup mirrors the newsvendor case except for the tuning grids: the shrinkage parameter λ ranges over a linear grid of 25 values in $[0, 1]$, the Wasserstein radius ρ over a logarithmic grid of 25 values in $[0.001, 5]$, and the Sinkhorn pair $(\bar{\rho}, \epsilon)$ over the 25 combinations $\bar{\rho} \in \{0.001, 0.005, 0.01, 0.05, 0.1\}$, $\epsilon \in \{0.01, 0.05, 0.1, 0.5, 1.0\}$. Training and test samples are drawn i.i.d. from the same short/long mixture as above, with $N = 10$ held fixed.

Figure 5 (left panel) reports the out-of-sample relative suboptimality over 100 random seeds, while Figure 5 (right panel) reports the average runtime, both as functions of n . The shrinkage path heuristic outperforms SAA, both regularized SAA benchmarks, and the SAA-RO endpoint heuristic at every n , while Wasserstein DRO and Sinkhorn DRO, in turn, outperform the shrinkage path heuristic. ℓ_1 -RSAA-CV and ℓ_2 -RSAA-CV are nearly indistinguishable from SAA at every n , indicating that cross-validation typically selects a vanishing regularization weight; shrinking the schedule toward $\mathbf{0}$ thus confers no robustness benefit here, whereas shrinking toward the RO solution does. However, Wasserstein and Sinkhorn DRO are three to four orders of magnitude slower than the shrinkage path heuristic across all n . The appointment scheduling experiment thus complements the newsvendor one: the shrinkage path heuristic yields competitive robust decisions that capture around 45%–70% of the median out-of-sample benefits of Wasserstein DRO over SAA, at a fraction of the computational cost. These percentages follow the same convention as for Figure 3.

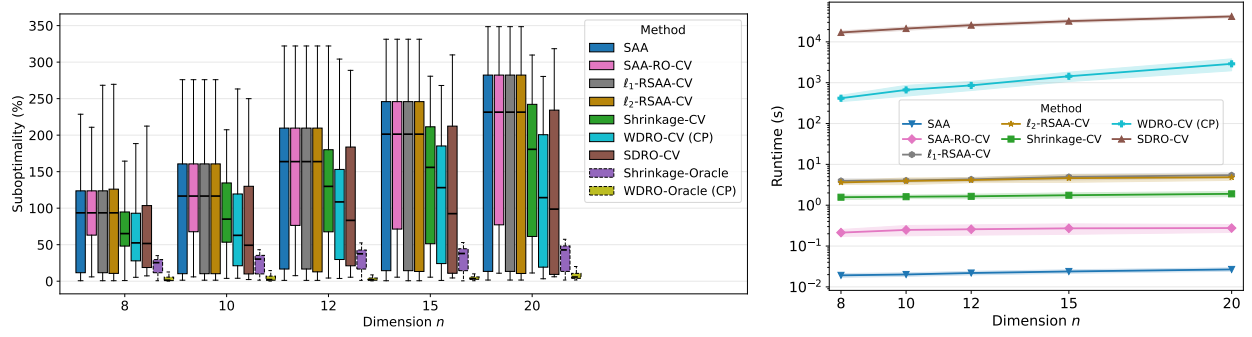


Figure 5 Out-of-sample suboptimality (left panel) and average runtime (right panel) of the different methods on the appointment scheduling problem with $N = 10$ samples, as the number of appointments varies over $n \in \{8, 10, 12, 15, 20\}$. The shrinkage path heuristic captures a significant part of the out-of-sample benefit of Wasserstein and Sinkhorn DRO over SAA, at orders of magnitude less computational cost, whereas the regularized SAA benchmarks (ℓ_1 -RSAA-CV and ℓ_2 -RSAA-CV) barely improve on SAA. Boxes, whiskers, and shaded bands follow the conventions of Figure 3.

Use of Large Language Models

Large language models (including Claude Fable 5 and Chat-GPT-Pro 5.6) were used to assist in writing the manuscript and implementing the numerical experiments. The authors have carefully reviewed and edited all generated content and take full responsibility for any remaining errors.

Endnotes

1. GitHub repository: <https://github.com/lingjun-meng/minimax-shrinkage>.

References

- Edward Anderson and Andy Philpott. Improving sample average approximation using distributional robustness. *INFORMS Journal on Optimization*, 4(1):90–124, 2022.
- Alvin J. Baranchik. A family of minimax estimators of the mean of a multivariate normal distribution. *The Annals of Mathematical Statistics*, 41(2):642–645, 1970.
- Amir Beck and Aharon Ben-Tal. Duality in robust optimization: Primal worst equals dual best. *Operations Research Letters*, 37(1):1–6, 2009.
- Aharon Ben-Tal, Laurent El Ghaoui, and Arkadi Nemirovski. *Robust Optimization*. Princeton Series in Applied Mathematics. Princeton University Press, Princeton, NJ, 2009.
- Amine Bennouna, Ryan Lucas, and Bart Van Parys. Certified robust neural networks: Generalization and corruption resistance. In *International Conference on Machine Learning*, pages 2092–2112. PMLR, 2023.
- Dimitris Bertsimas and Jack Dunn. *Machine Learning Under a Modern Optimization Lens*. Dynamic Ideas LLC, Charlestown, MA, 2019.

- Jose H. Blanchet and Karthyek R. A. Murthy. Quantifying distributional model risk via optimal transport. *Mathematics of Operations Research*, 44(2):565–600, 2019.
- Jose H. Blanchet, Karthyek R. A. Murthy, and Fan Zhang. Optimal transport-based distributionally robust optimization: Structural properties and iterative schemes. *Mathematics of Operations Research*, 47(2):1500–1529, 2022.
- Meysam Cheramin, Jianqiang Cheng, Ruiwei Jiang, and Kai Pan. Computationally efficient approximations for distributionally robust optimization under moment and Wasserstein ambiguity. *INFORMS Journal on Computing*, 34(3):1768–1794, 2022.
- Brian Denton and Diwakar Gupta. A sequential bounding approach for optimal appointment scheduling. *IIE Transactions*, 35(11):1003–1016, 2003.
- Adrián Esteban-Pérez and Juan M. Morales. Distributionally robust optimal power flow with contextual information. *European Journal of Operational Research*, 306(3):1047–1058, 2023.
- Rui Gao and Anton J. Kleywegt. Distributionally robust stochastic optimization with Wasserstein distance. *Mathematics of Operations Research*, 48(2):603–655, 2023.
- Rui Gao, Xi Chen, and Anton J. Kleywegt. Wasserstein distributionally robust optimization and variation regularization. *Operations Research*, 72(3):1177–1191, 2024.
- Grani A. Hanasusanto, Daniel Kuhn, Stein W. Wallace, and Steve Zymler. Distributionally robust multi-item newsvendor problems with multimodal demand distributions. *Mathematical Programming*, 152(1–2):1–32, 2015.
- Hao Hao and Peter Zhang. Robust paths: Geometry and computation. *arXiv preprint arXiv:2508.20039*, 2025.
- William James and Charles Stein. Estimation with quadratic loss. In *Proceedings of the Fourth Berkeley Symposium on Mathematical Statistics and Probability, Volume 1: Contributions to the Theory of Statistics*, pages 361–379, Berkeley, CA, 1961. University of California Press.
- Ruiwei Jiang, Minseok Ryu, and Guanglin Xu. Data-driven distributionally robust appointment scheduling over Wasserstein balls. *arXiv preprint arXiv:1907.03219*, 2019.
- Anton J. Kleywegt, Alexander Shapiro, and Tito Homem-de-Mello. The sample average approximation method for stochastic discrete optimization. *SIAM Journal on Optimization*, 12(2):479–502, 2002.
- Daniel Kuhn, Peyman Mohajerin Esfahani, Viet Anh Nguyen, and Soroosh Shafieezadeh-Abadeh. Wasserstein distributionally robust optimization: Theory and applications in machine learning. In *Operations Research & Management Science in the Age of Analytics*, INFORMS TutORials in Operations Research, pages 130–166. INFORMS, 2019.
- Daniel Kuhn, Soroosh Shafiee, and Wolfram Wiesemann. Distributionally robust optimization. *Acta Numerica*, 34:579–804, 2025.

- Jiajin Li, Sen Huang, and Anthony Man-Cho So. A first-order algorithmic framework for distributionally robust logistic regression. *Advances in Neural Information Processing Systems*, 32, 2019.
- Fengqiao Luo and Sanjay Mehrotra. Decomposition algorithm for distributionally robust optimization using Wasserstein metric with an application to a class of regression models. *European Journal of Operational Research*, 278(1):20–35, 2019.
- Peyman Mohajerin Esfahani and Daniel Kuhn. Data-driven distributionally robust optimization using the Wasserstein metric: Performance guarantees and tractable reformulations. *Mathematical Programming*, 171(1–2):115–166, 2018.
- Almir Mutapcic and Stephen Boyd. Cutting-set methods for robust convex optimization with pessimizing oracles. *Optimization Methods and Software*, 24(3):381–406, 2009.
- Panos M. Pardalos and Stephen A. Vavasis. Quadratic programming with one negative eigenvalue is NP-hard. *Journal of Global Optimization*, 1(1):15–22, 1991.
- Hamed Rahimian and Sanjay Mehrotra. Frameworks and results in distributionally robust optimization. *Open Journal of Mathematical Optimization*, 3:1–85, 2022. Article 4.
- Soroosh Shafiee, Liviu Aolaritei, Florian Dörfler, and Daniel Kuhn. Nash equilibria, regularization, and computation in optimal transport-based distributionally robust optimization. *Operations Research*, 74(3):1689–1709, 2026.
- Soroosh Shafieezadeh-Abadeh, Viet Anh Nguyen, Daniel Kuhn, and Peyman Mohajerin Esfahani. Wasserstein distributionally robust Kalman filtering. *Advances in Neural Information Processing Systems*, 31, 2018.
- Soroosh Shafieezadeh-Abadeh, Daniel Kuhn, and Peyman Mohajerin Esfahani. Regularization via mass transportation. *Journal of Machine Learning Research*, 20(103):1–68, 2019.
- Maurice Sion. On general minimax theorems. *Pacific Journal of Mathematics*, 8(1):171–176, 1958.
- Jie Wang, Rui Gao, and Yao Xie. Sinkhorn distributionally robust optimization. *Operations Research*, 74(3):1581–1603, 2026.
- Geoffrey S. Watson. Serial correlation in regression analysis. I. *Biometrika*, 42(3–4):327–341, 1955.
- Weijun Xie and Shabbir Ahmed. Distributionally robust chance constrained optimal power flow with renewables: A conic reformulation. *IEEE Transactions on Power Systems*, 33(2):1860–1867, 2018.
- Linwei Xin and David Alan Goldberg. Distributionally robust inventory control when demand is a martingale. *Mathematics of Operations Research*, 47(3):2387–2414, 2022.
- Jianzhe Zhen, Daniel Kuhn, and Wolfram Wiesemann. A unified theory of robust and distributionally robust optimization via the primal-worst-equals-dual-best principle. *Operations Research*, 73(2):862–878, 2025.
- Paul H. Zipkin. *Foundations of Inventory Management*. McGraw-Hill, Boston, 2000.

Supplementary Material

EC.1. Proofs of Main Theoretical Results

Proof of Theorem 1. We first show that the shrinkage path heuristic is exact whenever $\rho \geq \text{diam}(\Xi)$. This implies both the validity and the tightness of the stated bound for all ρ . We then establish the validity of the bound for $\rho \in [0, \text{diam}(\Xi)]$. Finally, we establish its pointwise tightness on the same interval by constructing, for each radius $\rho \in [0, \text{diam}(\Xi)]$, an instance which attains the bound for that ρ . Throughout the proof, we denote by f_ρ^* the optimal value of the WDRO problem (1), by $\hat{f}_\rho$ the objective value of the shrinkage path heuristic, and by $f_\rho(\mathbf{x})$ the objective value of (1) for a fixed decision $\mathbf{x}$.

First, observe that if $\text{diam}(\Xi) = 0$, then Ξ is a singleton and the heuristic is exact for each radius. Therefore, we assume for the remainder of the proof that $\text{diam}(\Xi) > 0$. Next, for $\rho \geq \text{diam}(\Xi)$, Assumption 1(i) and the definition of the Wasserstein distance then yield

$$W(\mathbb{Q}, \hat{\mathbb{P}}) = \inf_{\pi \in \Pi(\mathbb{Q}, \hat{\mathbb{P}})} \int_{\Xi \times \Xi} \|\xi - \xi'\|_2 \pi(d\xi, d\xi') \leq \sup_{\xi, \xi' \in \Xi} \|\xi - \xi'\|_2 = \text{diam}(\Xi).$$

Hence, when $\rho \geq \text{diam}(\Xi)$ the constraint $W(\mathbb{Q}, \hat{\mathbb{P}}) \leq \rho$ holds vacuously for every $\mathbb{Q} \in \mathcal{P}(\Xi)$, and $f_\rho(\mathbf{x}) = \sup_{\mathbb{Q} \in \mathcal{P}(\Xi)} \mathbb{E}_{\mathbb{Q}}[L(\tilde{\xi}^\top \mathbf{x})] = \sup_{\xi \in \Xi} L(\xi^\top \mathbf{x})$. The WDRO problem (1) thus reduces to the RO problem (3b), whose optimal value is attained by $\mathbf{x}(\infty)$. Since $\mathbf{x}(\infty)$ lies on the shrinkage path, the shrinkage path heuristic is exact. Likewise, our suboptimality bound yields $\rho \left[1 - \frac{\rho}{\text{diam}(\Xi)}\right]_+ \text{Lip}(L)\|\mathbf{x}(0)\|_2 = 0$ when $\rho \geq \text{diam}(\Xi)$, which therefore is both valid and tight.

We now prove the validity of our bound for $\rho \in [0, \text{diam}(\Xi)]$. To this end, we derive an upper bound on the heuristic's value $\hat{f}_\rho$ and a matching lower bound on the optimal value f_ρ^* .

As for the upper bound on the heuristic's value $\hat{f}_\rho$, we evaluate the objective $f_\rho(\mathbf{x})$ of (1) at the two extreme solutions $\mathbf{x}(0)$ and $\mathbf{x}(\infty)$ and then combine them to obtain a bound. While $\mathbf{x}(0)$ is optimal at $\rho = 0$ (i.e., $f_0(\mathbf{x}(0)) = f_0^*$), it is generally suboptimal once $\rho > 0$ (i.e., $f_\rho(\mathbf{x}(0)) \neq f_\rho^*$). To bound $f_\rho(\mathbf{x}(0))$, we control how fast the worst-case objective can grow as the radius increases to ρ . By Kantorovich–Rubinstein duality, $W(\mathbb{Q}, \hat{\mathbb{P}}) = \sup_{g: \text{Lip}(g) \leq 1} \{\mathbb{E}_{\mathbb{Q}}[g(\tilde{\xi})] - \mathbb{E}_{\hat{\mathbb{P}}}[g(\tilde{\xi})]\}$. Hence, every Lipschitz function g satisfies $\mathbb{E}_{\mathbb{Q}}[g(\tilde{\xi})] - \mathbb{E}_{\hat{\mathbb{P}}}[g(\tilde{\xi})] \leq \text{Lip}(g)W(\mathbb{Q}, \hat{\mathbb{P}})$. By Assumption 1(ii), the loss function L is Lipschitz continuous with constant $\text{Lip}(L)$, so for any $\mathbf{x}$ the map $\xi \mapsto L(\xi^\top \mathbf{x})$ is Lipschitz in ξ with constant $\text{Lip}(L)\|\mathbf{x}\|_2$. Applying the Kantorovich–Rubinstein inequality to $g(\xi) = L(\xi^\top \mathbf{x})$ and invoking $f_0(\mathbf{x}) = \mathbb{E}_{\hat{\mathbb{P}}}[L(\tilde{\xi}^\top \mathbf{x})]$, we obtain $\mathbb{E}_{\mathbb{Q}}[L(\tilde{\xi}^\top \mathbf{x})] \leq f_0(\mathbf{x}) + \text{Lip}(L)\|\mathbf{x}\|_2 W(\mathbb{Q}, \hat{\mathbb{P}})$. Taking the supremum over all $\mathbb{Q} \in \mathbb{B}_\rho(\hat{\mathbb{P}})$ with $W(\mathbb{Q}, \hat{\mathbb{P}}) \leq \rho$ yields $f_\rho(\mathbf{x}) \leq f_0(\mathbf{x}) + \rho \text{Lip}(L)\|\mathbf{x}\|_2$. Applying the preceding inequality at $\mathbf{x} = \mathbf{x}(0)$ and observing that $f_0(\mathbf{x}(0)) = f_0^*$ by construction, we have

$$f_\rho(\mathbf{x}(0)) \leq f_0(\mathbf{x}(0)) + \rho \text{Lip}(L)\|\mathbf{x}(0)\|_2 = f_0^* + \rho \text{Lip}(L)\|\mathbf{x}(0)\|_2. \quad (\text{EC.1})$$

To bound $f_\rho(\mathbf{x}(\infty))$, observe that for each fixed $\mathbf{x}$ the map $\rho \mapsto f_\rho(\mathbf{x})$ is nondecreasing, because the Wasserstein ball expands as ρ increases. Hence, for all $\rho \in [0, \text{diam}(\Xi)]$, we have

$$f_\rho(\mathbf{x}(\infty)) \leq f_{\text{diam}(\Xi)}(\mathbf{x}(\infty)) = f_{\text{diam}(\Xi)}^*, \quad (\text{EC.2})$$

where the final equality holds because $\mathbf{x}(\infty)$ optimally solves the WDRO problem (1) at $\rho = \text{diam}(\Xi)$. It remains to combine the bounds (EC.1) and (EC.2). Note that $f_\rho(\mathbf{x})$ is convex in $\mathbf{x}$ since it constitutes a pointwise supremum of weighted sums of the convex functions $\mathbf{x} \mapsto L(\boldsymbol{\xi}^\top \mathbf{x})$. We fix $\lambda = \frac{\rho}{\text{diam}(\Xi)} \in [0, 1]$ on the shrinkage path and apply Jensen's inequality to f_ρ to obtain

$$\begin{aligned} \hat{f}_\rho &\leq f_\rho\left(\left(1 - \frac{\rho}{\text{diam}(\Xi)}\right)\mathbf{x}(0) + \frac{\rho}{\text{diam}(\Xi)}\mathbf{x}(\infty)\right) \\ &\leq \left(1 - \frac{\rho}{\text{diam}(\Xi)}\right)f_\rho(\mathbf{x}(0)) + \frac{\rho}{\text{diam}(\Xi)}f_\rho(\mathbf{x}(\infty)) \\ &\leq \left(1 - \frac{\rho}{\text{diam}(\Xi)}\right)\left(f_0^* + \rho \text{Lip}(L)\|\mathbf{x}(0)\|_2\right) + \frac{\rho}{\text{diam}(\Xi)}f_{\text{diam}(\Xi)}^*, \end{aligned} \quad (\text{EC.3})$$

where the last inequality invokes (EC.1) and (EC.2).

As for the lower bound on the optimal value f_ρ^* , we exploit the concavity of f_ρ^* in ρ : a concave function lies above the chord through any two of its points, so f_ρ^* is bounded below on $[0, \text{diam}(\Xi)]$ by the chord joining its values at $\rho = 0$ and $\rho = \text{diam}(\Xi)$. To establish this concavity, we invoke the dual formulation of $f_\rho(\mathbf{x})$ (Mohajerin Esfahani and Kuhn 2018, Theorem 4.2), namely $f_\rho(\mathbf{x}) = \inf_{\eta \geq 0} \left\{ \eta \rho + \frac{1}{N} \sum_{i=1}^N \sup_{\boldsymbol{\xi} \in \Xi} [\ell(\mathbf{x}, \boldsymbol{\xi}) - \eta \|\boldsymbol{\xi} - \hat{\boldsymbol{\xi}}_i\|_2] \right\}$. As each $f_\rho(\mathbf{x})$ is an infimum of functions affine in ρ , the optimal value $f_\rho^* = \min_{\mathbf{x} \in \mathcal{X}(\Xi)} f_\rho(\mathbf{x})$ is itself a pointwise infimum of affine functions of ρ , and therefore concave in ρ . Writing any $\rho \in [0, \text{diam}(\Xi)]$ as the convex combination $\rho = \left(1 - \frac{\rho}{\text{diam}(\Xi)}\right)0 + \frac{\rho}{\text{diam}(\Xi)}\text{diam}(\Xi)$, concavity then yields

$$f_\rho^* \geq \left(1 - \frac{\rho}{\text{diam}(\Xi)}\right)f_0^* + \frac{\rho}{\text{diam}(\Xi)}f_{\text{diam}(\Xi)}^*. \quad (\text{EC.4})$$

Finally, subtracting the lower bound (EC.4) on f_ρ^* from the upper bound (EC.3) on $\hat{f}_\rho$, we obtain

$$\hat{f}_\rho - f_\rho^* \leq \left(1 - \frac{\rho}{\text{diam}(\Xi)}\right)\rho \text{Lip}(L)\|\mathbf{x}(0)\|_2.$$

The bound is parabolic in ρ and is thus maximized at $\rho = \frac{\text{diam}(\Xi)}{2}$.

We close by establishing the pointwise tightness of our bound for $\rho \in [0, \text{diam}(\Xi)]$. As observed above, $f_0^* = f_0(\mathbf{x}(0))$ and $f_{\text{diam}(\Xi)}^* = f_{\text{diam}(\Xi)}(\mathbf{x}(\infty))$, with both $\mathbf{x}(0)$ and $\mathbf{x}(\infty)$ lying on the shrinkage path. Hence, the heuristic is exact—and our bound tight—at both endpoints. For each remaining radius $\bar{\rho} \in (0, \text{diam}(\Xi))$, we construct an instance for which the bound is attained at $\rho = \bar{\rho}$. This instance has loss $L(z) = |z|$ (which satisfies Assumption 1(ii) with $\text{Lip}(L) = 1$), a single sample $\hat{\mathbb{P}} = \delta_{\boldsymbol{\xi}^0}$, and support $\Xi = \text{conv}\{\boldsymbol{\xi}^0, \boldsymbol{\xi}^1\}$, where $\boldsymbol{\xi}^0 = (0, \text{diam}(\Xi))^\top$ and $\boldsymbol{\xi}^1 = (\text{diam}(\Xi), \text{diam}(\Xi))^\top$. The

feasible set is $\mathcal{X}(\Xi) = \text{conv}\left\{(1, 0)^\top, \left(0, \frac{\bar{\rho}}{\text{diam}(\Xi)}\right)^\top, \left(\frac{\bar{\rho}}{\text{diam}(\Xi)}, 0\right)^\top\right\}$. On this instance, since $\mathbf{x} \geq \mathbf{0}$ and $\boldsymbol{\xi} \geq \mathbf{0}$, the loss is linear: $L(\boldsymbol{\xi}^\top \mathbf{x}) = \boldsymbol{\xi}^\top \mathbf{x}$. Because $\hat{\mathbb{P}}$ is supported on a single atom, the Wasserstein distance reduces to an expectation under $\mathbb{Q} \in \mathcal{P}(\Xi)$: $W(\mathbb{Q}, \hat{\mathbb{P}}) = \mathbb{E}_{\mathbb{Q}}[\|\tilde{\boldsymbol{\xi}} - \boldsymbol{\xi}^0\|_2] = \mathbb{E}_{\mathbb{Q}}[\tilde{\xi}_1]$, where the second equality follows from the support structure of this instance. Hence, the worst-case objective $f_\rho(\mathbf{x})$ admits the closed form

$$\begin{aligned} f_\rho(\mathbf{x}) &= \sup_{\mathbb{Q} \in \mathcal{P}(\Xi)} \left\{ \mathbb{E}_{\mathbb{Q}}[\tilde{\xi}_1 x_1 + \tilde{\xi}_2 x_2] : \mathbb{E}_{\mathbb{Q}}[\tilde{\xi}_1] \leq \rho \right\} \\ &= x_2 \text{diam}(\Xi) + x_1 \sup_{\mathbb{Q}_1 \in \mathcal{P}([0, \text{diam}(\Xi)])} \left\{ \mathbb{E}_{\mathbb{Q}_1}[\tilde{\xi}_1] : \mathbb{E}_{\mathbb{Q}_1}[\tilde{\xi}_1] \leq \rho \right\} \\ &= x_2 \text{diam}(\Xi) + x_1 \min\{\rho, \text{diam}(\Xi)\}. \end{aligned}$$

The two path endpoints are $\mathbf{x}(0) = (1, 0)^\top$, the minimizer of the SAA objective $f_0(\mathbf{x}) = \text{diam}(\Xi)x_2$ over $\mathcal{X}(\Xi)$, and $\mathbf{x}(\infty) = \left(0, \frac{\bar{\rho}}{\text{diam}(\Xi)}\right)^\top$, the minimizer of the RO objective $f_{\text{diam}(\Xi)}(\mathbf{x}) = \text{diam}(\Xi)(x_1 + x_2)$ over $\mathcal{X}(\Xi)$. Both endpoints satisfy $f_{\bar{\rho}}(\mathbf{x}(0)) = f_{\bar{\rho}}(\mathbf{x}(\infty)) = \bar{\rho}$, so by linearity $f_{\bar{\rho}}(\mathbf{x}) = \bar{\rho}$ for every $\mathbf{x}$ on the shrinkage path, giving $\hat{f}_{\bar{\rho}} = \bar{\rho}$. To evaluate the optimal value $f_\rho^* = \min_{\mathbf{x} \in \mathcal{X}(\Xi)} f_\rho(\mathbf{x})$, note that f_ρ is affine in $\mathbf{x}$ and hence attains its minimum at a vertex of $\mathcal{X}(\Xi)$. Evaluating f_ρ at the three vertices of $\mathcal{X}(\Xi)$, we find that $(1, 0)^\top$ and $\left(0, \frac{\bar{\rho}}{\text{diam}(\Xi)}\right)^\top$ both give $\bar{\rho}$, whereas $\left(\frac{\bar{\rho}}{\text{diam}(\Xi)}, 0\right)^\top$ gives $\frac{\bar{\rho}^2}{\text{diam}(\Xi)} < \bar{\rho}$ since $\bar{\rho} < \text{diam}(\Xi)$. Hence $f_\rho^* = \frac{\bar{\rho}^2}{\text{diam}(\Xi)}$, and therefore $\hat{f}_\rho - f_\rho^* = \bar{\rho} - \frac{\bar{\rho}^2}{\text{diam}(\Xi)} = \bar{\rho}\left(1 - \frac{\bar{\rho}}{\text{diam}(\Xi)}\right) \text{Lip}(L)\|\mathbf{x}(0)\|_2$. This is exactly the bound from the statement of the theorem, thereby establishing tightness. $\square$

Proof of Proposition 1. Throughout the proof, we use the same notation f_ρ^* , $\hat{f}_\rho$, and $f_\rho(\mathbf{x})$ as in the proof of Theorem 1. If $\bar{r} = 0$, then each sample is zero and the only radius in $[0, \bar{r}]$ at which the gap is zero is $\rho = 0$. We therefore assume throughout the proof that $\bar{r} > 0$. We first recall a closed-form identity for the WDRO objective value $f_\rho(\mathbf{x})$. Under Assumption 1 with $\Xi = \mathbb{R}^k$, the WDRO problem reduces to a regularized SAA problem (Shafieezadeh-Abadeh et al. 2019, Theorem 4(ii)), yielding, for every $\mathbf{x} \in \mathcal{X}(\Xi)$,

$$f_\rho(\mathbf{x}) = \frac{1}{N} \sum_{i=1}^N L(\hat{\boldsymbol{\xi}}_i^\top \mathbf{x}) + \rho \text{Lip}(L)\|\mathbf{x}\|_2. \quad (\text{EC.5})$$

When $\rho \rightarrow \infty$, the penalty on $\|\mathbf{x}\|_2$ dominates $f_\rho(\mathbf{x})$. The WDRO endpoint solution is therefore the minimum-norm solution $\mathbf{x}(\infty) = \arg \min_{\mathbf{x} \in \mathcal{X}(\Xi)} \|\mathbf{x}\|_2$, which exists and is unique since $\mathcal{X}(\Xi)$ is nonempty, closed and convex.

To prove the bound, we combine an upper bound on the heuristic's value $\hat{f}_\rho$, obtained from the identity (EC.5) at $\mathbf{x}(0)$ and $\mathbf{x}(\infty)$, with a lower bound on the optimal value f_ρ^* , obtained from the concavity of $\rho \mapsto f_\rho^*$ on $[0, \bar{r}]$. Subtracting the latter from the former bounds the suboptimality $\hat{f}_\rho - f_\rho^*$ up to a residual term proportional to $f_0(\mathbf{x}(\infty)) - f_{\bar{r}}^*$, which we show to be nonpositive. Together these steps give the claimed bound.

To bound the heuristic's value $\hat{f}_\rho$ from above, we evaluate (EC.5) at $\mathbf{x}(0)$ and $\mathbf{x}(\infty)$, which yields

$$f_\rho(\mathbf{x}(0)) = f_0^* + \rho \text{Lip}(L) \|\mathbf{x}(0)\|_2, \quad (\text{EC.6})$$

$$f_\rho(\mathbf{x}(\infty)) = f_0(\mathbf{x}(\infty)) + \rho \text{Lip}(L) \|\mathbf{x}(\infty)\|_2. \quad (\text{EC.7})$$

where the first identity additionally invokes the optimality of $\mathbf{x}(0)$ at $\rho = 0$.

Fixing $\lambda = \frac{\rho}{\bar{r}} \in [0, 1]$ on the shrinkage path and using the convexity of $f_\rho(\mathbf{x})$ in $\mathbf{x}$, Jensen's inequality yields an upper bound on the heuristic's value $\hat{f}_\rho$,

$$\begin{aligned} \hat{f}_\rho &\leq f_\rho\left(\left(1 - \frac{\rho}{\bar{r}}\right)\mathbf{x}(0) + \frac{\rho}{\bar{r}}\mathbf{x}(\infty)\right) \\ &\leq \left(1 - \frac{\rho}{\bar{r}}\right)f_\rho(\mathbf{x}(0)) + \frac{\rho}{\bar{r}}f_\rho(\mathbf{x}(\infty)) \\ &= \left(1 - \frac{\rho}{\bar{r}}\right)(f_0^* + \rho \text{Lip}(L) \|\mathbf{x}(0)\|_2) + \frac{\rho}{\bar{r}}(f_0(\mathbf{x}(\infty)) + \rho \text{Lip}(L) \|\mathbf{x}(\infty)\|_2), \end{aligned} \quad (\text{EC.8})$$

where the equality invokes (EC.6) and (EC.7).

To bound the optimal value f_ρ^* from below, we note that by the identity (EC.5), $f_\rho(\mathbf{x})$ is affine in ρ , and therefore $\rho \mapsto f_\rho^* = \min_{\mathbf{x} \in \mathcal{X}(\Xi)} f_\rho(\mathbf{x})$ is concave as the pointwise infimum of affine functions of ρ . Writing any $\rho \in [0, \bar{r}]$ as the convex combination $\rho = (1 - \frac{\rho}{\bar{r}})0 + \frac{\rho}{\bar{r}}\bar{r}$, the concavity of $\rho \mapsto f_\rho^*$ yields the following lower bound on f_ρ^* :

$$f_\rho^* \geq \left(1 - \frac{\rho}{\bar{r}}\right) f_0^* + \frac{\rho}{\bar{r}} f_{\bar{r}}^*. \quad (\text{EC.9})$$

Subtracting the lower bound (EC.9) from the upper bound (EC.8) yields

$$\begin{aligned} \hat{f}_\rho - f_\rho^* &\leq \left(1 - \frac{\rho}{\bar{r}}\right) (f_0^* + \rho \text{Lip}(L) \|\mathbf{x}(0)\|_2) + \frac{\rho}{\bar{r}} (f_0(\mathbf{x}(\infty)) + \rho \text{Lip}(L) \|\mathbf{x}(\infty)\|_2) - \left(1 - \frac{\rho}{\bar{r}}\right) f_0^* - \frac{\rho}{\bar{r}} f_{\bar{r}}^* \\ &= \left(1 - \frac{\rho}{\bar{r}}\right) \rho \text{Lip}(L) \|\mathbf{x}(0)\|_2 + \frac{\rho^2}{\bar{r}} \text{Lip}(L) \|\mathbf{x}(\infty)\|_2 + \frac{\rho}{\bar{r}} (f_0(\mathbf{x}(\infty)) - f_{\bar{r}}^*). \end{aligned}$$

It remains to show that $f_0(\mathbf{x}(\infty)) \leq f_{\bar{r}}^*$ to yield the claimed bound in Proposition 1. Observe that $f_0(\mathbf{x}) = \frac{1}{N} \sum_{i=1}^N L(\hat{\boldsymbol{\xi}}_i^\top \mathbf{x})$. By the Lipschitz continuity of L and the Cauchy–Schwarz inequality, for all $\mathbf{x} \in \mathcal{X}(\Xi)$, we have

$$\begin{aligned} f_0(\mathbf{x}(\infty)) &\leq f_0(\mathbf{x}) + \frac{1}{N} \sum_{i=1}^N \text{Lip}(L) |\hat{\boldsymbol{\xi}}_i^\top (\mathbf{x} - \mathbf{x}(\infty))| \\ &\leq f_0(\mathbf{x}) + \left(\frac{1}{N} \sum_{i=1}^N \|\hat{\boldsymbol{\xi}}_i\|_2\right) \text{Lip}(L) \|\mathbf{x} - \mathbf{x}(\infty)\|_2 \\ &= f_0(\mathbf{x}) + \bar{r} \text{Lip}(L) \|\mathbf{x} - \mathbf{x}(\infty)\|_2. \end{aligned} \quad (\text{EC.10})$$

Since $\mathbf{x}(\infty)$ minimizes $\|\cdot\|_2$ —equivalently $\frac{1}{2}\|\cdot\|_2^2$ —over the convex set $\mathcal{X}(\Xi)$, the first-order optimality condition $\nabla\left(\frac{1}{2}\|\mathbf{x}\|_2^2\right)_{\mathbf{x}(\infty)}^\top (\mathbf{x} - \mathbf{x}(\infty)) \geq 0$ gives $\mathbf{x}(\infty)^\top (\mathbf{x} - \mathbf{x}(\infty)) \geq 0$ for all $\mathbf{x} \in \mathcal{X}(\Xi)$. Therefore,

we have $\|\mathbf{x}\|_2^2 = \|\mathbf{x} - \mathbf{x}(\infty)\|_2^2 + 2\mathbf{x}(\infty)^\top(\mathbf{x} - \mathbf{x}(\infty)) + \|\mathbf{x}(\infty)\|_2^2 \geq \|\mathbf{x} - \mathbf{x}(\infty)\|_2^2$. Substituting $\|\mathbf{x} - \mathbf{x}(\infty)\|_2 \leq \|\mathbf{x}\|_2$ into (EC.10) and invoking the identity (EC.5) at $\rho = \bar{r}$, we have

$$f_0(\mathbf{x}(\infty)) \leq f_0(\mathbf{x}) + \bar{r}\text{Lip}(L)\|\mathbf{x}\|_2 = f_{\bar{r}}(\mathbf{x}).$$

The inequality $f_0(\mathbf{x}(\infty)) \leq f_{\bar{r}}(\mathbf{x})$ holds for all $\mathbf{x} \in \mathcal{X}(\Xi)$. Minimizing the right-hand side over $\mathbf{x} \in \mathcal{X}(\Xi)$ yields the target inequality $f_0(\mathbf{x}(\infty)) \leq f_{\bar{r}}^*$, and therefore

$$\begin{aligned} \hat{f}_\rho - f_\rho^* &\leq \left(1 - \frac{\rho}{\bar{r}}\right) \rho \text{Lip}(L)\|\mathbf{x}(0)\|_2 + \frac{\rho^2}{\bar{r}} \text{Lip}(L)\|\mathbf{x}(\infty)\|_2 + \frac{\rho}{\bar{r}} (f_0(\mathbf{x}(\infty)) - f_{\bar{r}}^*) \\ &\leq \text{Lip}(L) \left(\rho \left(1 - \frac{\rho}{\bar{r}}\right) \|\mathbf{x}(0)\|_2 + \frac{\rho^2}{\bar{r}} \|\mathbf{x}(\infty)\|_2 \right), \end{aligned}$$

which completes the proof of the claimed bound. For $\mathbf{x}(\infty) = \mathbf{0}$, this bound is maximized at $\rho = \bar{r}/2$.

□

The proof of Theorem 2 relies on the following auxiliary results, which we state and prove first. Throughout, we use the notation f_ρ^* , $\hat{f}_\rho$, and $f_\rho(\mathbf{x})$ from the proof of Theorem 1, we write $\Delta(\rho) := \hat{f}_\rho - f_\rho^*$ for the suboptimality gap of the shrinkage path heuristic at the radius ρ , and we denote by $\mathbf{x}^*(\rho)$ and $\hat{\mathbf{x}}(\rho)$ the corresponding optimizers of the WDRO problem (1) and of the shrinkage path heuristic, respectively. The first lemma expresses the optimal values f_ρ^* and $\hat{f}_\rho$ through explicit quadratics; the second lemma locates a gap-maximizing radius ρ_{gap} and expresses $\Delta(\rho_{\text{gap}})$ as the covariance between two functions of a random eigenvalue of $\mathbf{H}$; the third lemma bounds this covariance by an *endpoint envelope*, that is, an upper bound that depends on the spectrum of $\mathbf{H}$ only through its extreme eigenvalues $\lambda_{\min}$ and $\lambda_{\max}$; and the fourth lemma maximizes the resulting univariate envelope in closed form.

LEMMA EC.1. *Let Assumption 2 hold and suppose $\mathbf{x}(0) \neq \mathbf{0}$. After absorbing $\text{Lip}(L)$ into ρ , so that $f_\rho(\mathbf{x}) = f_0(\mathbf{x}) + \rho\|\mathbf{x}\|_2$, the WDRO optimal value satisfies, for every $\rho \geq 0$,*

$$f_\rho^* = f_0^* + \min_{\mathbf{x} \in \mathbb{R}^n} \{(\mathbf{x} - \mathbf{x}(0))^\top \mathbf{H}(\mathbf{x} - \mathbf{x}(0)) + \rho\|\mathbf{x}\|_2\}. \quad (\text{EC.11})$$

Moreover, $\mathbf{x}(\infty) = \mathbf{0}$, and the shrinkage path solution satisfies $\hat{\mathbf{x}}(\rho) = \lambda^*(\rho)\mathbf{x}(0)$, where

$$\lambda^*(\rho) = \left[1 - \frac{\rho\|\mathbf{x}(0)\|_2}{2\mathbf{x}(0)^\top \mathbf{H} \mathbf{x}(0)} \right]_+, \quad (\text{EC.12})$$

and the corresponding shrinkage path value is

$$\hat{f}_\rho = f_0^* + \mathbf{x}(0)^\top \mathbf{H} \mathbf{x}(0) - \mathbf{x}(0)^\top \mathbf{H} \mathbf{x}(0) \left[1 - \frac{\rho\|\mathbf{x}(0)\|_2}{2\mathbf{x}(0)^\top \mathbf{H} \mathbf{x}(0)} \right]_+^2. \quad (\text{EC.13})$$

Proof of Lemma EC.1. The normalization of the radius is harmless: under the present assumptions, $\mathcal{Z}$ contains a nondegenerate interval and L is a nonconstant quadratic on it, so $\text{Lip}(L) > 0$ and the substitution $\rho \leftarrow \text{Lip}(L)\rho$ simply relabels the radii $\rho \geq 0$. The remainder of the proof proceeds in three steps: the first step shows that on the ball $B := \{\mathbf{x} \in \mathbb{R}^n : \|\mathbf{x}\|_2 \leq \|\mathbf{x}(0)\|_2\}$, the empirical loss f_0 is a quadratic centered at $\mathbf{x}(0)$; the second step shows that the WDRO problem and the quadratic problem in (EC.11) both attain their minima on B , where their objectives coincide up to the constant f_0^* , which establishes (EC.11); and the third step determines $\mathbf{x}(\infty)$ and optimizes over the shrinkage path, which yields (EC.12) and (EC.13).

For the first step, the definition of $\mathcal{Z}$ ensures that $\hat{\boldsymbol{\xi}}_i^\top \mathbf{x} \in \mathcal{Z}$ for all $\mathbf{x} \in B$ and $i \in [N]$, so Assumption 2(ii) implies that

$$f_0(\mathbf{x}) = \mathbf{x}^\top \mathbf{H} \mathbf{x} + b \mathbf{x}^\top \left(\frac{1}{N} \sum_{i=1}^N \hat{\boldsymbol{\xi}}_i \right) + c \quad \text{for all } \mathbf{x} \in B.$$

Because L is differentiable on $\mathbb{R}$ and agrees with the quadratic $z \mapsto z^2 + bz + c$ on $\mathcal{Z}$, we have $L'(z) = 2z + b$ for all $z \in \mathcal{Z}$: at interior points this is immediate, and at the endpoints $\pm R$ of $\mathcal{Z} := [-R, R]$, where $R := \|\mathbf{x}(0)\|_2 \max_{i \in [N]} \|\hat{\boldsymbol{\xi}}_i\|_2$, differentiability forces L' to coincide with the one-sided derivative of the quadratic taken from inside $\mathcal{Z}$. Since $\hat{\boldsymbol{\xi}}_i^\top \mathbf{x}(0) \in \mathcal{Z}$ for all i and $\mathcal{X}(\Xi) = \mathbb{R}^n$, the SAA first-order condition $\nabla f_0(\mathbf{x}(0)) = \mathbf{0}$ is

$$2\mathbf{H}\mathbf{x}(0) + b \left(\frac{1}{N} \sum_{i=1}^N \hat{\boldsymbol{\xi}}_i \right) = \mathbf{0}.$$

Substituting $b \frac{1}{N} \sum_{i=1}^N \hat{\boldsymbol{\xi}}_i = -2\mathbf{H}\mathbf{x}(0)$ into the expression for f_0 above and completing the square, we obtain

$$f_0(\mathbf{x}) = (\mathbf{x} - \mathbf{x}(0))^\top \mathbf{H} (\mathbf{x} - \mathbf{x}(0)) + c - \mathbf{x}(0)^\top \mathbf{H} \mathbf{x}(0) = f_0^* + (\mathbf{x} - \mathbf{x}(0))^\top \mathbf{H} (\mathbf{x} - \mathbf{x}(0)) \quad \text{for all } \mathbf{x} \in B,$$

where setting $\mathbf{x} = \mathbf{x}(0)$ in the first equality of the preceding display identifies the constant $c - \mathbf{x}(0)^\top \mathbf{H} \mathbf{x}(0)$ as $f_0(\mathbf{x}(0)) = f_0^*$.

For the second step, denote by $q(\mathbf{x}) := (\mathbf{x} - \mathbf{x}(0))^\top \mathbf{H} (\mathbf{x} - \mathbf{x}(0)) + \rho \|\mathbf{x}\|_2$ the function minimized in (EC.11). In view of the first step and the identity (EC.5), we have $f_\rho(\mathbf{x}) = f_0(\mathbf{x}) + \rho \|\mathbf{x}\|_2 = f_0^* + q(\mathbf{x})$ for all $\mathbf{x} \in B$. Both f_ρ and q attain their minima over $\mathbb{R}^n$ on the compact set B : for $\rho = 0$, both functions are minimized by $\mathbf{x}(0) \in B$, while for $\rho > 0$, every $\mathbf{x} \notin B$ satisfies $\|\mathbf{x}\|_2 > \|\mathbf{x}(0)\|_2$ and is therefore dominated by $\mathbf{x}(0)$, since the optimality of $\mathbf{x}(0)$ in the SAA problem (3a) implies

$$f_\rho(\mathbf{x}) = f_0(\mathbf{x}) + \rho \|\mathbf{x}\|_2 > f_0(\mathbf{x}(0)) + \rho \|\mathbf{x}(0)\|_2 = f_\rho(\mathbf{x}(0)) \quad \text{and} \quad q(\mathbf{x}) \geq \rho \|\mathbf{x}\|_2 > \rho \|\mathbf{x}(0)\|_2 = q(\mathbf{x}(0)).$$

We thus conclude that

$$f_\rho^* = \min_{\mathbf{x} \in B} f_\rho(\mathbf{x}) = f_0^* + \min_{\mathbf{x} \in B} q(\mathbf{x}) = f_0^* + \min_{\mathbf{x} \in \mathbb{R}^n} q(\mathbf{x}),$$

which is (EC.11); in particular, the minimizers of the WDRO problem coincide with those of q .

For the third step, $\mathcal{X}(\Xi) = \mathbb{R}^n$ implies that the minimum-norm endpoint is $\mathbf{x}(\infty) = \mathbf{0}$, so the shrinkage path is $\{\lambda \mathbf{x}(0) : \lambda \in [0, 1]\} \subseteq B$, where λ denotes the coefficient of $\mathbf{x}(0)$ and thus corresponds to $1 - \lambda$ in the parameterization (4). The first step and the identity (EC.5) then give

$$f_\rho(\lambda \mathbf{x}(0)) = f_0^* + (1 - \lambda)^2 \mathbf{x}(0)^\top \mathbf{H} \mathbf{x}(0) + \rho \lambda \|\mathbf{x}(0)\|_2.$$

Minimizing this strongly convex univariate quadratic over $\lambda \in [0, 1]$ yields the minimizer $\lambda^*(\rho)$ in (EC.12), so that $\hat{\mathbf{x}}(\rho) = \lambda^*(\rho) \mathbf{x}(0)$; substituting this minimizer into the path objective gives (EC.13). $\square$

LEMMA EC.2. *Let the conditions of Lemma EC.1 hold, and let ρ_{gap} be a maximizer of $\Delta(\rho)$ over $\rho \geq 0$ with $\Delta(\rho_{\text{gap}}) > 0$. Then $\rho_{\text{gap}} \in (0, \rho_0^h)$, where*

$$\rho_0^h := \frac{2 \mathbf{x}(0)^\top \mathbf{H} \mathbf{x}(0)}{\|\mathbf{x}(0)\|_2}.$$

Moreover,

$$\|\hat{\mathbf{x}}(\rho_{\text{gap}})\|_2 = \|\mathbf{x}^*(\rho_{\text{gap}})\|_2. \quad (\text{EC.14})$$

If $\mathbf{H} = \mathbf{U} \text{diag}(\lambda_1, \dots, \lambda_n) \mathbf{U}^\top$ with $\mathbf{U}$ orthogonal, $\tilde{\mathbf{x}}(0) = \mathbf{U}^\top \mathbf{x}(0)$, $p_i = \tilde{x}(0)_i^2 / \|\mathbf{x}(0)\|_2^2$, and $\phi_\gamma(\lambda) = (\lambda / (\lambda + \gamma))^2$, then, with $\mathbb{P}(\Lambda = \lambda_i) = p_i$ and $\gamma_{\text{gap}} = \rho_{\text{gap}} / (2 \|\mathbf{x}^*(\rho_{\text{gap}})\|_2)$,

$$\frac{\Delta(\rho_{\text{gap}})}{\|\mathbf{x}(0)\|_2^2} = \mathbb{E}_p[\Lambda \phi_{\gamma_{\text{gap}}}(\Lambda)] - \mathbb{E}_p[\Lambda] \mathbb{E}_p[\phi_{\gamma_{\text{gap}}}(\Lambda)], \quad (\text{EC.15})$$

and

$$\mathbb{E}_p[\phi_{\gamma_{\text{gap}}}(\Lambda)] = \left(\frac{\mathbb{E}_p[\Lambda]}{\mathbb{E}_p[\Lambda] + \gamma_{\text{gap}}} \right)^2 = \phi_{\gamma_{\text{gap}}}(\mathbb{E}_p[\Lambda]). \quad (\text{EC.16})$$

Proof of Lemma EC.2. We proceed in three steps: first, we establish that $\rho_{\text{gap}} \in (0, \rho_0^h)$ and the norm identity (EC.14), second, we derive closed forms for $f_{\rho_{\text{gap}}}^*$ and $\hat{f}_{\rho_{\text{gap}}}$ from the first-order condition of the quadratic problem in (EC.11), which yield the covariance identity (EC.15) in the eigenbasis of $\mathbf{H}$, and third, we rewrite (EC.14) in this eigenbasis to obtain (EC.16).

For the first step, since $\Delta(0) = 0 < \Delta(\rho_{\text{gap}})$, the maximizer ρ_{gap} satisfies $\rho_{\text{gap}} > 0$. The objective in (EC.11) is strongly convex, and the path objective is strongly convex in λ , so both $\mathbf{x}^*(\rho)$ and $\hat{\mathbf{x}}(\rho)$ are unique. Danskin's theorem states that a minimum-value function $\rho \mapsto \min_u g(u, \rho)$, whose objective is affine in ρ and whose minimum is attained at a unique $u(\rho)$, is differentiable with derivative $(\partial g / \partial \rho)(u(\rho), \rho)$. Since ρ enters both minimization problems only through the term $\rho \|\mathbf{x}\|_2$, which is affine in ρ , we obtain, for every $\rho > 0$,

$$\Delta'(\rho) = \|\hat{\mathbf{x}}(\rho)\|_2 - \|\mathbf{x}^*(\rho)\|_2.$$

Fermat's rule, which states that the derivative of a differentiable function vanishes at every interior extremum, applied at the maximizer $\rho_{\text{gap}} > 0$, gives $\Delta'(\rho_{\text{gap}}) = 0$ and hence we obtain (EC.14). This rules out every $\rho_{\text{gap}} \geq \rho_0^h$: at any such radius (EC.12) gives $\hat{\mathbf{x}}(\rho_{\text{gap}}) = \mathbf{0}$, so (EC.14) would force $\mathbf{x}^*(\rho_{\text{gap}}) = \mathbf{0}$, and then (EC.11) and (EC.13) would both evaluate to $f_0^* + \mathbf{x}(0)^\top \mathbf{H} \mathbf{x}(0)$, making the heuristic exact and contradicting $\Delta(\rho_{\text{gap}}) > 0$. Thus $\rho_{\text{gap}} \in (0, \rho_0^h)$; moreover, $\hat{\mathbf{x}}(\rho_{\text{gap}}) = \lambda^*(\rho_{\text{gap}}) \mathbf{x}(0) \neq \mathbf{0}$, and so $\mathbf{x}^*(\rho_{\text{gap}}) \neq \mathbf{0}$ by (EC.14).

For the second step, recall from Lemma EC.1 that $\mathbf{x}^*(\rho_{\text{gap}})$ minimizes the quadratic problem on the right-hand side of (EC.11); the associated first-order condition at $\mathbf{x}^*(\rho_{\text{gap}}) \neq \mathbf{0}$ is

$$2\mathbf{H}(\mathbf{x}^*(\rho_{\text{gap}}) - \mathbf{x}(0)) + \rho_{\text{gap}} \frac{\mathbf{x}^*(\rho_{\text{gap}})}{\|\mathbf{x}^*(\rho_{\text{gap}})\|_2} = \mathbf{0}. \quad (\text{EC.17})$$

With $\gamma_{\text{gap}} := \rho_{\text{gap}} / (2\|\mathbf{x}^*(\rho_{\text{gap}})\|_2)$, this is equivalently

$$(\mathbf{H} + \gamma_{\text{gap}} \mathbf{I}) \mathbf{x}^*(\rho_{\text{gap}}) = \mathbf{H} \mathbf{x}(0).$$

Multiplying this equation from the left by $\mathbf{x}^*(\rho_{\text{gap}})^\top$ yields $\mathbf{x}(0)^\top \mathbf{H} \mathbf{x}^*(\rho_{\text{gap}}) = \mathbf{x}^*(\rho_{\text{gap}})^\top \mathbf{H} \mathbf{x}^*(\rho_{\text{gap}}) + \gamma_{\text{gap}} \|\mathbf{x}^*(\rho_{\text{gap}})\|_2^2$; substituting this identity and $\rho_{\text{gap}} \|\mathbf{x}^*(\rho_{\text{gap}})\|_2 = 2\gamma_{\text{gap}} \|\mathbf{x}^*(\rho_{\text{gap}})\|_2^2$ into the objective of (EC.11), evaluated at $\mathbf{x}^*(\rho_{\text{gap}})$, yields

$$f_{\rho_{\text{gap}}}^* = f_0^* + \mathbf{x}(0)^\top \mathbf{H} \mathbf{x}(0) - \mathbf{x}^*(\rho_{\text{gap}})^\top \mathbf{H} \mathbf{x}^*(\rho_{\text{gap}}).$$

Similarly, since $\hat{\mathbf{x}}(\rho_{\text{gap}}) = \lambda^*(\rho_{\text{gap}}) \mathbf{x}(0)$ implies $\hat{\mathbf{x}}(\rho_{\text{gap}})^\top \mathbf{H} \hat{\mathbf{x}}(\rho_{\text{gap}}) = \lambda^*(\rho_{\text{gap}})^2 \mathbf{x}(0)^\top \mathbf{H} \mathbf{x}(0)$, the path value (EC.13) can be written as

$$\hat{f}_{\rho_{\text{gap}}} = f_0^* + \mathbf{x}(0)^\top \mathbf{H} \mathbf{x}(0) - \hat{\mathbf{x}}(\rho_{\text{gap}})^\top \mathbf{H} \hat{\mathbf{x}}(\rho_{\text{gap}}).$$

Subtracting the first identity from the second and using $\hat{\mathbf{x}}(\rho_{\text{gap}})^\top \mathbf{H} \hat{\mathbf{x}}(\rho_{\text{gap}}) = \frac{\|\hat{\mathbf{x}}(\rho_{\text{gap}})\|_2^2}{\|\mathbf{x}(0)\|_2^2} \mathbf{x}(0)^\top \mathbf{H} \mathbf{x}(0)$ together with the norm identity (EC.14) gives

$$\Delta(\rho_{\text{gap}}) = \mathbf{x}^*(\rho_{\text{gap}})^\top \mathbf{H} \mathbf{x}^*(\rho_{\text{gap}}) - \frac{\|\mathbf{x}^*(\rho_{\text{gap}})\|_2^2}{\|\mathbf{x}(0)\|_2^2} \mathbf{x}(0)^\top \mathbf{H} \mathbf{x}(0). \quad (\text{EC.18})$$

Diagonalize $\mathbf{H} = \mathbf{U} \text{diag}(\lambda_1, \dots, \lambda_n) \mathbf{U}^\top$ with $\mathbf{U}$ orthogonal, set $\tilde{\mathbf{x}}(0) = \mathbf{U}^\top \mathbf{x}(0)$, and define $p_i := \tilde{\mathbf{x}}(0)_i^2 / \|\mathbf{x}(0)\|_2^2$. Then $p_i \geq 0$ and $\sum_i p_i = 1$. Writing the resolvent equation $(\mathbf{H} + \gamma_{\text{gap}} \mathbf{I}) \mathbf{x}^*(\rho_{\text{gap}}) = \mathbf{H} \mathbf{x}(0)$ in this eigenbasis gives

$$\tilde{x}_i^*(\rho_{\text{gap}}) = \frac{\lambda_i}{\lambda_i + \gamma_{\text{gap}}} \tilde{x}_i(0).$$

Substituting this expression into (EC.18) yields

$$\Delta(\rho_{\text{gap}}) = \|\mathbf{x}(0)\|_2^2 \left[\sum_{i=1}^n \lambda_i \left(\frac{\lambda_i}{\lambda_i + \gamma_{\text{gap}}} \right)^2 p_i - \sum_{i=1}^n \left(\frac{\lambda_i}{\lambda_i + \gamma_{\text{gap}}} \right)^2 p_i \sum_{j=1}^n \lambda_j p_j \right]. \quad (\text{EC.19})$$

Define the random variable Λ by $\mathbb{P}(\Lambda = \lambda_i) = p_i$, and set $\phi_\gamma(\lambda) := (\lambda/(\lambda + \gamma))^2$. Then (EC.19) becomes the covariance identity (EC.15). Because $\phi_{\gamma_{\text{gap}}}$ is increasing in λ , this covariance is nonnegative; the gap is precisely the cost of applying one common shrinkage factor instead of the coordinatewise factors $\lambda_i/(\lambda_i + \gamma_{\text{gap}})$.

For the third step, it remains to rewrite the norm identity (EC.14) in spectral form. Let $m := \mathbb{E}_p[\Lambda]$ and $\lambda_F := \lambda^*(\rho_{\text{gap}})$. From the eigenbasis representation,

$$\|\mathbf{x}^*(\rho_{\text{gap}})\|_2^2 = \|\mathbf{x}(0)\|_2^2 \mathbb{E}_p[\phi_{\gamma_{\text{gap}}}(\Lambda)],$$

and (EC.14) gives $\mathbb{E}_p[\phi_{\gamma_{\text{gap}}}(\Lambda)] = \lambda_F^2$. On the other hand, since $\rho_{\text{gap}} \in (0, \rho_0^h)$,

$$\lambda_F = 1 - \frac{\rho_{\text{gap}} \|\mathbf{x}(0)\|_2}{2 \mathbf{x}(0)^\top \mathbf{H} \mathbf{x}(0)}.$$

Using $\rho_{\text{gap}} = 2\gamma_{\text{gap}}\lambda_F\|\mathbf{x}(0)\|_2$ —which follows from the definition of γ_{gap} together with $\|\mathbf{x}^*(\rho_{\text{gap}})\|_2 = \lambda_F\|\mathbf{x}(0)\|_2$, that is, (EC.14)—and $m = \mathbf{x}(0)^\top \mathbf{H} \mathbf{x}(0)/\|\mathbf{x}(0)\|_2^2$, this becomes $\lambda_F = 1 - \gamma_{\text{gap}}\lambda_F/m$, or

$$\lambda_F = \frac{m}{m + \gamma_{\text{gap}}}.$$

Consequently, (EC.16) holds. $\square$

LEMMA EC.3. *Under the conditions of Lemma EC.2, the gap at the maximizing radius ρ_{gap} admits a spectral upper bound that depends on $\mathbf{H}$ only through its condition number κ . Specifically,*

$$\frac{\Delta(\rho_{\text{gap}})}{\lambda_{\min}\|\mathbf{x}(0)\|_2^2} \leq \sup_{t \geq 0} t^2 \left(\frac{1}{\sqrt{1+t}} - \frac{1}{\sqrt{\kappa+t}} \right)^2. \quad (\text{EC.20})$$

Proof of Lemma EC.3. The proof proceeds in three steps: the first step rewrites the covariance in (EC.15) as the difference between $\mathbb{E}_p[r_\gamma(\Lambda)]$ and $r_\gamma(\mathbb{E}_p[\Lambda])$, that is, the Jensen gap of the resolvent $r_\gamma(\lambda) := 1/(\lambda + \gamma)$; the second step bounds this gap by relaxing the law of Λ to a mean-constrained distribution on the spectral interval $[\lambda_{\min}, \lambda_{\max}]$; and the third step maximizes the resulting endpoint expression over the mean and rescales it to obtain (EC.20).

For the first step, write γ for γ_{gap} , set $m := \mathbb{E}_p[\Lambda]$, and define $r_\gamma(\lambda) := 1/(\lambda + \gamma)$. Using the covariance identity (EC.15), the spectral stationarity condition (EC.16) in the form $\mathbb{E}_p[\phi_\gamma(\Lambda)] = \phi_\gamma(m)$, and the identity

$$(\lambda + \gamma)\phi_\gamma(\lambda) = \lambda - \gamma + \gamma^2 r_\gamma(\lambda),$$

we obtain the resolvent representation

$$\frac{\Delta(\rho_{\text{gap}})}{\|\mathbf{x}(0)\|_2^2} = \mathbb{E}_p[\Lambda\phi_\gamma(\Lambda)] - m\phi_\gamma(m) = \gamma^2(\mathbb{E}_p[r_\gamma(\Lambda)] - r_\gamma(m)). \quad (\text{EC.21})$$

Indeed, the middle expression of (EC.21) equals

$$m - \gamma + \gamma^2 \mathbb{E}_p[r_\gamma(\Lambda)] - (m + \gamma)\phi_\gamma(m),$$

and $(m + \gamma)\phi_\gamma(m) = m - \gamma + \gamma^2 r_\gamma(m)$. Thus the stationarity condition (EC.16) has reduced the covariance to a resolvent Jensen gap. From this point on, the argument uses only the support interval $[\lambda_{\min}, \lambda_{\max}]$ and the mean m of Λ .

For the second step, we relax the requirements on the law of Λ and merely require it to be supported on $[\lambda_{\min}, \lambda_{\max}]$ with mean $\mathbb{E}_p[\Lambda] = m$. Since r_γ is convex on $[\lambda_{\min}, \lambda_{\max}]$, every such law satisfies the secant bound

$$\mathbb{E}_p[r_\gamma(\Lambda)] \leq \frac{\lambda_{\max} - m}{\lambda_{\max} - \lambda_{\min}} r_\gamma(\lambda_{\min}) + \frac{m - \lambda_{\min}}{\lambda_{\max} - \lambda_{\min}} r_\gamma(\lambda_{\max}). \quad (\text{EC.22})$$

The right-hand side of (EC.22) equals $\mathbb{E}_p[r_\gamma(\Lambda)]$ when Λ follows the two-point law on $\{\lambda_{\min}, \lambda_{\max}\}$ with mean m . The secant bound therefore holds with equality for this law. Substituting (EC.22) into (EC.21) gives the endpoint envelope

$$\frac{\Delta(\rho_{\text{gap}})}{\|\mathbf{x}(0)\|_2^2} \leq \frac{\gamma^2(m - \lambda_{\min})(\lambda_{\max} - m)}{(\lambda_{\min} + \gamma)(\lambda_{\max} + \gamma)(m + \gamma)}.$$

The resulting expression is nevertheless an upper bound, not an exact reduction of the stationarity-constrained problem, because stationarity is no longer imposed after the resolvent rewrite. Taking the supremum over $m \in [\lambda_{\min}, \lambda_{\max}]$ and $\gamma \geq 0$ yields

$$\frac{\Delta(\rho_{\text{gap}})}{\|\mathbf{x}(0)\|_2^2} \leq \sup_{m \in [\lambda_{\min}, \lambda_{\max}], \gamma \geq 0} \frac{\gamma^2(m - \lambda_{\min})(\lambda_{\max} - m)}{(\lambda_{\min} + \gamma)(\lambda_{\max} + \gamma)(m + \gamma)}. \quad (\text{EC.23})$$

For the third step, fix $\gamma > 0$. The only m -dependent factor in (EC.23) is $(m - \lambda_{\min})(\lambda_{\max} - m)/(m + \gamma)$. Its derivative is

$$\frac{(\lambda_{\min} + \gamma)(\lambda_{\max} + \gamma) - (m + \gamma)^2}{(m + \gamma)^2},$$

so the maximizing mean is

$$m^*(\gamma) = \sqrt{(\lambda_{\min} + \gamma)(\lambda_{\max} + \gamma)} - \gamma,$$

which lies in $[\lambda_{\min}, \lambda_{\max}]$. Evaluating (EC.23) at $m^*(\gamma)$ reduces the bound to

$$\frac{\Delta(\rho_{\text{gap}})}{\|\mathbf{x}(0)\|_2^2} \leq \sup_{\gamma \geq 0} \gamma^2 \left(\frac{1}{\sqrt{\lambda_{\min} + \gamma}} - \frac{1}{\sqrt{\lambda_{\max} + \gamma}} \right)^2.$$

With $\gamma = \lambda_{\min} t$ and $\kappa = \lambda_{\max}/\lambda_{\min}$, this is (EC.20). □

LEMMA EC.4. *The scalar envelope in Lemma EC.3 admits the following closed-form bound: for every $\kappa \geq 1$ and $t \geq 0$,*

$$t^2 \left(\frac{1}{\sqrt{1+t}} - \frac{1}{\sqrt{\kappa+t}} \right)^2 \leq \frac{(\kappa-1)^2}{\varphi^5(\kappa-1) + 27}. \quad (\text{EC.24})$$

Proof of Lemma EC.4. The claim is immediate when $\kappa = 1$, so take $\kappa > 1$. The proof proceeds in two steps: the first step rationalizes the difference of square roots, then rescales the desired inequality through variables u and z ; the second step maximizes the resulting left-hand side over u and compares the maximum with the right-hand side.

For the first step, we invoke the identity $(\sqrt{\kappa+t} - \sqrt{1+t})(\sqrt{\kappa+t} + \sqrt{1+t}) = \kappa - 1$ to rewrite the left-hand side of (EC.24) as $t^2(\kappa-1)^2 / [(1+t)(\kappa+t)(\sqrt{1+t} + \sqrt{\kappa+t})^2]$. Clearing the positive denominators and dividing by $(\kappa-1)^2 > 0$, we obtain the equivalent inequality

$$(\varphi^5(\kappa-1) + 27) t^2 \leq (1+t)(\kappa+t) (\sqrt{1+t} + \sqrt{\kappa+t})^2.$$

Now set $u := t/(1+t) \in [0, 1]$ and $z := (\kappa-1)/(1+t) \geq 0$, and define $R(z) := (1+z)(1+\sqrt{1+z})^2$. Dividing the preceding inequality by $(1+t)^3$ shows that it suffices to prove, for all $u \in [0, 1]$ and $z \geq 0$,

$$u^2 [27(1-u) + \varphi^5 z] \leq R(z). \quad (\text{EC.25})$$

For the second step, fix z . The derivative of the left-hand side with respect to u is $u(54 + 2\varphi^5 z - 81u)$. Hence, with $z_0 := 27/(2\varphi^5)$,

$$\max_{u \in [0, 1]} u^2 [27(1-u) + \varphi^5 z] = \begin{cases} 4(1 + \varphi^5 z/27)^3, & 0 \leq z \leq z_0, \\ \varphi^5 z, & z \geq z_0, \end{cases}$$

where the first branch is attained at the interior critical point and the second at $u = 1$. We now compare both branches with $R(z)$. For the second branch, writing $s = \sqrt{1+z} \geq 1$ gives

$$R(z) - \varphi^5 z = (s+1)(s-\varphi)^2(s+2+\sqrt{5}) \geq 0.$$

For the first branch, set $g(z) := (R(z)/4)^{1/3}$. We claim that g is concave. Indeed, writing $s = \sqrt{1+z}$ gives $g(z) = 2^{-2/3}(s(s+1))^{2/3}$ and

$$g''(z) = -\frac{2^{-2/3}(4s^2 + 7s + 4)}{18s^{10/3}(s+1)^{4/3}} < 0.$$

Moreover $g(0) = 1$, and the preceding branch bound at z_0 gives $g(z_0) \geq (\varphi^5 z_0/4)^{1/3} = 3/2$. Therefore, for $z \in [0, z_0]$,

$$g(z) \geq \left(1 - \frac{z}{z_0}\right) g(0) + \frac{z}{z_0} g(z_0) \geq 1 + \frac{\varphi^5 z}{27}.$$

Cubing both sides and invoking $g(z) = (R(z)/4)^{1/3}$ yields $R(z)/4 \geq (1 + \varphi^5 z/27)^3$, hence $R(z) \geq 4(1 + \varphi^5 z/27)^3$ on $[0, z_0]$. The two branch bounds $R(z) \geq 4(1 + \varphi^5 z/27)^3$ on $[0, z_0]$ and $R(z) \geq \varphi^5 z$

on $[z_0, \infty)$ together give $\max_{u \in [0,1]} u^2 [27(1-u) + \varphi^5 z] \leq R(z)$ for every $z \geq 0$, which proves (EC.25) and hence (EC.24). $\square$

Proof of Theorem 2. We use the notation introduced before Lemma EC.1. The proof proceeds in three steps: the first step handles the trivial case $\mathbf{x}(0) = \mathbf{0}$ and otherwise combines the auxiliary lemmas to prove the additive bound; the second step builds two-dimensional instances and stationary radii at which the gap can be evaluated explicitly; and the third step evaluates these stationary radii in the regimes $\kappa \rightarrow \infty$ and $\kappa \rightarrow 1^+$ to prove asymptotic tightness.

For the first step, if $\mathbf{x}(0) = \mathbf{0}$, then $\mathbf{0}$ minimizes both the empirical loss and the norm regularizer in (EC.5), so it is a WDRO optimizer for every $\rho \geq 0$; the shrinkage path is also identically zero, and hence $\Delta(\rho) = 0$ for every ρ . We therefore assume $\mathbf{x}(0) \neq \mathbf{0}$ in the remainder of the proof and, as in Lemma EC.1, absorb the constant $\text{Lip}(L)$ in (EC.5) into the radius ρ , which leaves the worst-case suboptimality $\sup_{\rho \geq 0} \Delta(\rho)$ unchanged.

Define $\rho_0^h := 2\mathbf{x}(0)^\top \mathbf{H} \mathbf{x}(0) / \|\mathbf{x}(0)\|_2$, and let ρ_{gap} maximize $\Delta(\rho)$ over $\rho \geq 0$. Such a maximizer exists because Lemma EC.1 gives continuity of both $\hat{f}_\rho$ and f_ρ^* in ρ , while $\hat{f}_\rho$ is constant for $\rho \geq \rho_0^h$ by (EC.13) and f_ρ^* is nondecreasing in ρ , so that Δ is nonincreasing on $[\rho_0^h, \infty)$ and $\sup_{\rho \geq 0} \Delta(\rho)$ is attained on the compact interval $[0, \rho_0^h]$. If $\Delta(\rho_{\text{gap}}) = 0$, the theorem is immediate. Otherwise, Lemmas EC.2, EC.3, and EC.4 give

$$\Delta(\rho_{\text{gap}}) \leq \lambda_{\min} \|\mathbf{x}(0)\|_2^2 \frac{(\kappa - 1)^2}{\varphi^5(\kappa - 1) + 27} = \lambda_{\max} \|\mathbf{x}(0)\|_2^2 \frac{(\kappa - 1)^2}{\kappa[\varphi^5(\kappa - 1) + 27]},$$

which is the bound (5).

For the second step, we set up the tightness construction. Consider the matrix $\mathbf{H}_\kappa$ and the SAA solution $\mathbf{x}_\kappa(0)$ given by

$$\mathbf{H}_\kappa = \text{diag}(1, \kappa), \quad \mathbf{x}_\kappa(0) = \|\mathbf{x}_\kappa(0)\|_2 (\sqrt{1-\theta}, \sqrt{\theta})^\top, \quad \theta \in (0, 1), \quad (\text{EC.26})$$

with spectral law $p_\theta = (1-\theta)\delta_1 + \theta\delta_\kappa$ and mean $m = 1 + \theta(\kappa - 1)$, where δ_a denotes the unit point mass at a . These reduced data—in fact, any $\mathbf{H} \succ \mathbf{0}$ and any target $\mathbf{x}(0) \neq \mathbf{0}$ —are realized by empirical samples satisfying Assumption 2. Fix $\beta > 2\sqrt{\mathbf{x}(0)^\top \mathbf{H} \mathbf{x}(0)}$ and set $\bar{\boldsymbol{\xi}} := -2\mathbf{H} \mathbf{x}(0)/\beta$. Then $\bar{\boldsymbol{\xi}}^\top \mathbf{H}^{-1} \bar{\boldsymbol{\xi}} = 4\mathbf{x}(0)^\top \mathbf{H} \mathbf{x}(0)/\beta^2 < 1$, so $\mathbf{H} - \bar{\boldsymbol{\xi}} \bar{\boldsymbol{\xi}}^\top \succ \mathbf{0}$ admits a factorization $\mathbf{H} - \bar{\boldsymbol{\xi}} \bar{\boldsymbol{\xi}}^\top = \mathbf{L} \mathbf{L}^\top$. The $N = 2n$ equally weighted samples $\bar{\boldsymbol{\xi}} \pm \sqrt{n} \mathbf{L} \mathbf{e}_j$, $j \in [n]$, have empirical mean $\bar{\boldsymbol{\xi}}$ and empirical second-moment matrix $\bar{\boldsymbol{\xi}} \bar{\boldsymbol{\xi}}^\top + \mathbf{L} \mathbf{L}^\top = \mathbf{H}$. Choose for L a Huber-type loss that agrees with $z^2 + \beta z + c$ on an open interval containing $\mathcal{Z}$ and continues linearly, with matching value and slope, outside that interval. Here c is large enough to ensure $L \geq 0$. This loss is convex, globally Lipschitz, differentiable, and quadratic on $\mathcal{Z}$, so Assumptions 1 and 2(ii) hold with $b = \beta$. Finally, the choice of $\bar{\boldsymbol{\xi}}$ yields the first-order condition $\nabla f_0(\mathbf{x}(0)) = 2\mathbf{H} \mathbf{x}(0) + \beta \bar{\boldsymbol{\xi}} = \mathbf{0}$, and f_0 coincides with the strongly convex

quadratic $\mathbf{x} \mapsto \mathbf{x}^\top \mathbf{H} \mathbf{x} + \beta \bar{\boldsymbol{\xi}}^\top \mathbf{x} + c$ on a neighborhood of $\mathbf{x}(0)$. Since a convex function with a strong local minimizer admits no other minimizers, $\mathbf{x}(0)$ is the unique SAA solution.

To evaluate the gaps produced by this construction, we use the same covariance derivation as in Lemma EC.2: the derivation only relies on the first-order condition (EC.17) and the stationarity property (EC.14), and is therefore not specific to the maximizing radius ρ_{gap} . We now construct radii with these two properties. For any $\gamma > 0$, the point $\mathbf{x}_\gamma := (\mathbf{H}_\kappa + \gamma \mathbf{I})^{-1} \mathbf{H}_\kappa \mathbf{x}_\kappa(0)$ is nonzero because $\mathbf{H}_\kappa \succ \mathbf{0}$ and $\mathbf{x}_\kappa(0) \neq \mathbf{0}$. It satisfies the first-order condition (EC.17) with $\rho(\gamma) := 2\gamma \|\mathbf{x}_\gamma\|_2$ in place of ρ_{gap} , hence $\mathbf{x}^*(\rho(\gamma)) = \mathbf{x}_\gamma$ by strong convexity of (EC.11). The dual multiplier $\rho(\gamma)/(2\|\mathbf{x}^*(\rho(\gamma))\|_2)$ equals γ itself. If, in addition, γ satisfies the spectral stationarity condition (EC.16) for the law p_θ with mean $m = \mathbb{E}_{p_\theta}[\Lambda]$ —in which case we call (θ, γ) a *stationary pair* and $\rho(\gamma)$ its *stationary radius*—then $\|\mathbf{x}_\gamma\|_2 = \|\mathbf{x}_\kappa(0)\|_2 \sqrt{\mathbb{E}_{p_\theta}[\phi_\gamma(\Lambda)]} = m(m + \gamma)^{-1} \|\mathbf{x}_\kappa(0)\|_2$, so that $\rho(\gamma) = 2\gamma m(m + \gamma)^{-1} \|\mathbf{x}_\kappa(0)\|_2$. Moreover, since $\rho(\gamma) \|\mathbf{x}_\kappa(0)\|_2 / (2\mathbf{x}_\kappa(0)^\top \mathbf{H}_\kappa \mathbf{x}_\kappa(0)) = \gamma / (m + \gamma) < 1$, the positive part in (EC.12) is inactive and $\lambda^*(\rho(\gamma)) = m / (m + \gamma)$. Thus $\|\hat{\mathbf{x}}(\rho(\gamma))\|_2 = \|\mathbf{x}_\gamma\|_2$, and the stationarity property (EC.14) indeed holds at $\rho(\gamma)$. For the two-point law p_θ , the covariance formula at a stationary radius therefore reduces to

$$\frac{\Delta(\rho(\gamma))}{\|\mathbf{x}_\kappa(0)\|_2^2} = \mathbb{E}_{p_\theta}[\Lambda \phi_\gamma(\Lambda)] - \mathbb{E}_{p_\theta}[\Lambda] \mathbb{E}_{p_\theta}[\phi_\gamma(\Lambda)] = \theta(1 - \theta)(\kappa - 1) [\phi_\gamma(\kappa) - \phi_\gamma(1)]. \quad (\text{EC.27})$$

Since $\lambda_{\max} = \kappa$, the upper bound (5) becomes $\|\mathbf{x}_\kappa(0)\|_2^2 (\kappa - 1)^2 / [\varphi^5 (\kappa - 1) + 27]$ on these instances. Thus it suffices to exhibit stationary radii whose gaps match this expression asymptotically.

For the third step, first let $\kappa \rightarrow \infty$ and set $\alpha := 1/\varphi$, so $\alpha^2 + \alpha = 1$ and $1 + \alpha = 1/\alpha = \varphi$. Take $\gamma_\kappa = \alpha\kappa$ and choose θ_κ to satisfy (EC.16), which for the law p_θ reads $(1 - \theta) \phi_{\alpha\kappa}(1) + \theta \phi_{\alpha\kappa}(\kappa) = \phi_{\alpha\kappa}(m)$ with $m = 1 + \theta(\kappa - 1)$. Since $\phi_{\alpha\kappa}(\kappa) = (1 + \alpha)^{-2} = \alpha^2$ for every κ , while $\phi_{\alpha\kappa}(1) \rightarrow 0$ and $\phi_{\alpha\kappa}(m) \rightarrow \theta^2 / (\theta + \alpha)^2$ as $\kappa \rightarrow \infty$, the stationarity equation converges to

$$\alpha^2 \theta = \frac{\theta^2}{(\theta + \alpha)^2},$$

whose roots in $(0, 1]$ are the interior point $\theta = \alpha^2$ —by the defining relation $\alpha^2 + \alpha = 1$ —and the boundary point $\theta = 1$. Define the limiting stationarity defect $D(\theta) := \alpha^2 \theta - \theta^2 / (\theta + \alpha)^2$. Its derivative at the interior root is $D'(\alpha^2) = \alpha^2(1 - 2\alpha) \neq 0$, so D changes sign at α^2 . The intermediate value theorem furnishes, for all sufficiently large κ , a root θ_κ of the exact stationarity equation with $\theta_\kappa \rightarrow \alpha^2$ and $1 - \theta_\kappa \rightarrow \alpha$. Let $\bar{\rho}_\kappa := \rho(\gamma_\kappa) = 2\gamma_\kappa m_\kappa (m_\kappa + \gamma_\kappa)^{-1} \|\mathbf{x}_\kappa(0)\|_2$ be the stationary radius of the pair $(\theta_\kappa, \gamma_\kappa)$, where $m_\kappa = 1 + \theta_\kappa(\kappa - 1)$. Since $\phi_{\alpha\kappa}(\kappa) = \alpha^2$ and $\phi_{\alpha\kappa}(1) \rightarrow 0$, (EC.27) gives

$$\frac{\Delta(\bar{\rho}_\kappa)}{\kappa \|\mathbf{x}_\kappa(0)\|_2^2} \longrightarrow \alpha^2 \alpha \alpha^2 = \frac{1}{\varphi^5}.$$

The upper bound divided by $\kappa \|\mathbf{x}_\kappa(0)\|_2^2$ has the same limit, proving tightness as $\kappa \rightarrow \infty$.

Finally let $\kappa = 1 + \epsilon$ with $\epsilon \downarrow 0$ and choose $\theta = 1/2$. The stationarity defect

$$D_\epsilon(\gamma) := \mathbb{E}_{p_{1/2}}[\phi_\gamma(\Lambda)] - \phi_\gamma\left(1 + \frac{\epsilon}{2}\right) = \frac{\epsilon^2}{8} \frac{2\gamma(\gamma-2)}{(1+\gamma)^4} + O(\epsilon^3)$$

vanishes at $\gamma = 2$ to leading order in ϵ , and the factor $2\gamma(\gamma-2)/(1+\gamma)^4$ has derivative $4/81 \neq 0$ at $\gamma = 2$. For all sufficiently small ϵ , the defect therefore changes sign near $\gamma = 2$, and the intermediate value theorem furnishes a stationary dual $\bar{\gamma}_\epsilon \rightarrow 2$. Let $\bar{\rho}_\epsilon := \rho(\bar{\gamma}_\epsilon)$ be the corresponding stationary radius. Since $\phi'_2(1) = 4/27$, (EC.27) gives

$$\frac{\Delta(\bar{\rho}_\epsilon)}{\|\mathbf{x}_\kappa(0)\|_2^2} = \frac{1}{4} \epsilon [\phi_{\bar{\gamma}_\epsilon}(1 + \epsilon) - \phi_{\bar{\gamma}_\epsilon}(1)] = \frac{\epsilon^2}{27} + O(\epsilon^3),$$

which matches $\epsilon^2/[\varphi^5\epsilon + 27] = \epsilon^2/27 + O(\epsilon^3)$. This proves tightness as $\kappa \rightarrow 1^+$ and completes the proof. $\square$

Proof of Corollary 1. For each κ , we use the two-point instance (EC.26) with $\theta = \theta_\kappa$, where θ_κ is the parameter constructed in the $\kappa \rightarrow \infty$ tightness step of the proof of Theorem 2 and satisfies $\theta_\kappa \rightarrow \alpha^2$ for $\alpha := 1/\varphi$. The proof proceeds in three steps: the first step shows that the maximal gap of the κ th instance asymptotically attains the worst-case bound (5); the second step identifies the asymptotic scale of the dual multiplier $\gamma_{\text{gap},\kappa}$ at a gap-maximizing radius; and the third step converts that scale into the three radii in the statement.

For the first step, let $\rho_{\text{gap},\kappa}$ be a radius that maximizes the gap $\Delta(\rho)$ for the κ th instance. The stationary radius $\bar{\rho}_\kappa$ constructed in the proof of Theorem 2 is an admissible candidate radius for the same instance, so $\Delta(\rho_{\text{gap},\kappa}) \geq \Delta(\bar{\rho}_\kappa)$. The lower bound from that construction and the upper bound (5) therefore squeeze the maximizing gap:

$$\frac{\Delta(\rho_{\text{gap},\kappa})}{\kappa \|\mathbf{x}_\kappa(0)\|_2^2} \longrightarrow \alpha^5 = \frac{1}{\varphi^5}. \quad (\text{EC.28})$$

For all sufficiently large κ , this gap is positive. Lemma EC.2 then applies at $\rho_{\text{gap},\kappa}$. We write $\gamma_{\text{gap},\kappa} := \rho_{\text{gap},\kappa}/(2\|\mathbf{x}_\kappa^*(\rho_{\text{gap},\kappa})\|_2)$ for the dual multiplier γ_{gap} from that lemma, specialized to the κ th instance.

For the second step, we claim that $\gamma_{\text{gap},\kappa}/\kappa \rightarrow \alpha$. For the two-point instances, $\lambda_{\min} = 1$ and $\lambda_{\max} = \kappa$. Write the scalar endpoint envelope in Lemma EC.3 as $G_\kappa(t) := t^2((1+t)^{-1/2} - (\kappa+t)^{-1/2})^2$. The proof of that lemma gives $\Delta(\rho_{\text{gap},\kappa})/\|\mathbf{x}_\kappa(0)\|_2^2 \leq G_\kappa(\gamma_{\text{gap},\kappa})$. Together with Lemma EC.4 and (EC.28), this implies the squeeze

$$\alpha^5 \longleftarrow \frac{\Delta(\rho_{\text{gap},\kappa})}{\kappa \|\mathbf{x}_\kappa(0)\|_2^2} \leq \frac{G_\kappa(\gamma_{\text{gap},\kappa})}{\kappa} \leq \frac{(\kappa-1)^2}{\kappa[\varphi^5(\kappa-1) + 27]} \longrightarrow \alpha^5. \quad (\text{EC.29})$$

Now set $\beta_\kappa := \gamma_{\text{gap},\kappa}/\kappa$ and define

$$h_\kappa(\beta) := \frac{G_\kappa(\beta\kappa)}{\kappa} = \beta^2 \left(\frac{1}{\sqrt{\beta+1/\kappa}} - \frac{1}{\sqrt{1+\beta}} \right)^2, \quad w(\beta) := \beta^2 \left(\frac{1}{\sqrt{\beta}} - \frac{1}{\sqrt{1+\beta}} \right)^2.$$

Since $1/\sqrt{\beta+1/\kappa} \rightarrow 1/\sqrt{\beta}$ uniformly on every compact subset of $(0, \infty)$, we have that $h_\kappa \rightarrow w$ uniformly on every such subset. Also, the conjugate form of G_κ and the inequalities $(\sqrt{1+t} + \sqrt{\kappa+t})^2 \geq \kappa+t$ and $(\kappa-1)^2 \leq \kappa^2$ yield

$$h_\kappa(\beta) \leq \frac{\beta^2 \kappa}{(1+\beta\kappa)(1+\beta)^2} \leq \frac{\beta}{(1+\beta)^2} \quad \forall \beta > 0.$$

Since the right-hand side tends to zero as $\beta \downarrow 0$ and as $\beta \uparrow \infty$, (EC.29) rules out $\beta_\kappa \rightarrow 0$ and $\beta_\kappa \rightarrow \infty$ along any subsequence. Along any subsequence with $\beta_\kappa \rightarrow \bar{\beta} \in (0, \infty)$, the local uniform convergence gives $w(\bar{\beta}) = \lim_{\kappa \rightarrow \infty} h_\kappa(\beta_\kappa) = \alpha^5$.

It remains only to identify the maximizer of w . Since $\sqrt{w(\beta)} = \sqrt{\beta} - \beta/\sqrt{1+\beta}$, the first-order condition for an interior maximizer of w is $(1+\beta)^3 = (2+\beta)^2\beta$, equivalently $\beta^2 + \beta - 1 = 0$. Its unique positive solution is $\beta = \alpha$. Also $w(\beta) \rightarrow 0$ as $\beta \downarrow 0$ and as $\beta \uparrow \infty$, so α is the unique maximizer of w , with $w(\alpha) = \alpha^5$. Hence every convergent subsequence of $\{\beta_\kappa\}$ has limit α , and the claim follows.

For the third step, let $m_\kappa := \mathbf{x}_\kappa(0)^\top \mathbf{H}_\kappa \mathbf{x}_\kappa(0) / \|\mathbf{x}_\kappa(0)\|_2^2 = 1 + \theta_\kappa(\kappa-1)$. Since $\theta_\kappa \rightarrow \alpha^2$, we have $m_\kappa/\kappa \rightarrow \alpha^2$. The stationarity relation in Lemma EC.2 and the limit $\gamma_{\text{gap},\kappa}/\kappa \rightarrow \alpha$ give

$$\lambda^*(\rho_{\text{gap},\kappa}) = \frac{m_\kappa}{m_\kappa + \gamma_{\text{gap},\kappa}} \longrightarrow \frac{\alpha^2}{\alpha^2 + \alpha} = \alpha^2.$$

The three characteristic radii are

$$\begin{aligned} \rho_{\text{gap},\kappa} &= 2\gamma_{\text{gap},\kappa} \lambda^*(\rho_{\text{gap},\kappa}) \|\mathbf{x}_\kappa(0)\|_2, & \rho_{0,\kappa}^h &= 2m_\kappa \|\mathbf{x}_\kappa(0)\|_2, \\ \rho_{0,\kappa}^* &= 2\|\mathbf{H}_\kappa \mathbf{x}_\kappa(0)\|_2 = 2\|\mathbf{x}_\kappa(0)\|_2 \sqrt{(1-\theta_\kappa) + \theta_\kappa \kappa^2}. \end{aligned}$$

The first identity combines the definition of $\gamma_{\text{gap},\kappa}$ with the norm identity (EC.14) and $\hat{\mathbf{x}}_\kappa(\rho) = \lambda^*(\rho) \mathbf{x}_\kappa(0)$; the second is the radius at which the shrinkage factor (EC.12) reaches zero; and the third holds because $\mathbf{x}_\kappa^*(\rho) = \mathbf{0}$ if and only if $\mathbf{0}$ satisfies the subgradient optimality condition $\rho \partial \|\cdot\|_2(\mathbf{0}) \ni 2\mathbf{H}_\kappa \mathbf{x}_\kappa(0)$ in the quadratic problem of (EC.11), that is, if and only if $\rho \geq 2\|\mathbf{H}_\kappa \mathbf{x}_\kappa(0)\|_2$. After dividing by $2\kappa\|\mathbf{x}_\kappa(0)\|_2$, these identities yield

$$\frac{\rho_{\text{gap},\kappa}}{2\kappa\|\mathbf{x}_\kappa(0)\|_2} \rightarrow \alpha^3, \quad \frac{\rho_{0,\kappa}^h}{2\kappa\|\mathbf{x}_\kappa(0)\|_2} \rightarrow \alpha^2, \quad \frac{\rho_{0,\kappa}^*}{2\kappa\|\mathbf{x}_\kappa(0)\|_2} \rightarrow \alpha.$$

Since $1/\alpha = \varphi$, we conclude that

$$\rho_{\text{gap},\kappa} : \rho_{0,\kappa}^h : \rho_{0,\kappa}^* \longrightarrow \alpha^3 : \alpha^2 : \alpha = 1 : \varphi : \varphi^2.$$

□

The proof of Theorem 3 relies on the following auxiliary results, which we state and prove first. Throughout, we use the notation introduced before Lemma EC.1 and the spectral notation of Lemma EC.2, and we write $\Delta^{0,\infty}(\rho) := \min\{f_\rho(\mathbf{x}(0)), f_\rho(\mathbf{x}(\infty))\} - f_\rho^*$ for the additive suboptimality

of the SAA–RO heuristic at the radius ρ . The first lemma locates ρ_{gap} , a maximizer of the suboptimality $\Delta(\rho)$ of the shrinkage path heuristic. The second lemma provides a correlation inequality in the spirit of Cassels’ classical reverse Cauchy–Schwarz inequality (Watson 1955, Appendix 1).

LEMMA EC.5. *Let the conditions of Lemma EC.2 hold. Then the gap-maximizing radius satisfies $\rho_{\text{gap}} > \rho_{\text{end}}$, where $\rho_{\text{end}} := \mathbf{x}(0)^\top \mathbf{H} \mathbf{x}(0) / \|\mathbf{x}(0)\|_2$.*

Proof of Lemma EC.5. Set $m := \mathbb{E}_p[\Lambda]$. Then $m = \mathbf{x}(0)^\top \mathbf{H} \mathbf{x}(0) / \|\mathbf{x}(0)\|_2^2 = \rho_{\text{end}} / \|\mathbf{x}(0)\|_2$. By Lemmas EC.1 and EC.2, $\|\mathbf{x}^*(\rho_{\text{gap}})\|_2 = \|\hat{\mathbf{x}}(\rho_{\text{gap}})\|_2 = \lambda^*(\rho_{\text{gap}}) \|\mathbf{x}(0)\|_2$ with $\lambda^*(\rho_{\text{gap}}) = m / (m + \gamma_{\text{gap}})$. Hence

$$\rho_{\text{gap}} = 2\gamma_{\text{gap}} \lambda^*(\rho_{\text{gap}}) \|\mathbf{x}(0)\|_2, \quad \rho_{\text{end}} = m \|\mathbf{x}(0)\|_2, \quad \frac{\rho_{\text{gap}}}{\rho_{\text{end}}} = \frac{2\gamma_{\text{gap}}}{m + \gamma_{\text{gap}}}.$$

It is therefore enough to show that $\gamma_{\text{gap}} > m$.

Suppose, toward a contradiction, that $\gamma_{\text{gap}} \leq m$. Let $a(\lambda) := \phi_{\gamma_{\text{gap}}}(m) + \phi'_{\gamma_{\text{gap}}}(m)(\lambda - m)$ be the tangent line to $\phi_{\gamma_{\text{gap}}}$ at m , and set $h := a - \phi_{\gamma_{\text{gap}}}$. Then $h(m) = h'(m) = 0$ and

$$h''(\lambda) = -\phi''_{\gamma_{\text{gap}}}(\lambda) = \frac{2\gamma_{\text{gap}}(2\lambda - \gamma_{\text{gap}})}{(\lambda + \gamma_{\text{gap}})^4}.$$

Thus h is concave on $(0, \gamma_{\text{gap}}/2)$ and strictly convex on $(\gamma_{\text{gap}}/2, \infty)$. Since $m \geq \gamma_{\text{gap}}$, convexity and $h(m) = h'(m) = 0$ imply $h \geq 0$ on $[\gamma_{\text{gap}}/2, \infty)$, with equality only at m . Moreover,

$$h(0) = \phi_{\gamma_{\text{gap}}}(m) - \phi'_{\gamma_{\text{gap}}}(m)m - \phi_{\gamma_{\text{gap}}}(0) = \frac{m^2}{(m + \gamma_{\text{gap}})^2} - \frac{2\gamma_{\text{gap}}m^2}{(m + \gamma_{\text{gap}})^3} = \frac{m^2(m - \gamma_{\text{gap}})}{(m + \gamma_{\text{gap}})^3} \geq 0.$$

Since also $h(\gamma_{\text{gap}}/2) > 0$ (were it zero, convexity and $h \geq 0$ on $[\gamma_{\text{gap}}/2, m]$ together with $h(m) = 0$ would force $h = 0$ on this interval, contradicting strict convexity), concavity on $[0, \gamma_{\text{gap}}/2]$ gives

$$h(\lambda) \geq \frac{\gamma_{\text{gap}}/2 - \lambda}{\gamma_{\text{gap}}/2} h(0) + \frac{\lambda}{\gamma_{\text{gap}}/2} h(\gamma_{\text{gap}}/2) > 0, \quad \lambda \in (0, \gamma_{\text{gap}}/2].$$

Thus $h(\lambda) \geq 0$ for all $\lambda > 0$, with equality only at $\lambda = m$. Taking expectations and using the stationarity condition (EC.16) gives

$$\mathbb{E}_p[h(\Lambda)] = \mathbb{E}_p[a(\Lambda)] - \mathbb{E}_p[\phi_{\gamma_{\text{gap}}}(\Lambda)] = a(m) - \phi_{\gamma_{\text{gap}}}(m) = 0.$$

Hence $\Lambda = m$ p -almost surely. The covariance formula (EC.15) would then give $\Delta(\rho_{\text{gap}}) = 0$, contradicting $\Delta(\rho_{\text{gap}}) > 0$. Therefore $\gamma_{\text{gap}} > m$, and so $\rho_{\text{gap}} > \rho_{\text{end}}$. $\square$

LEMMA EC.6. *Let X and Y be random variables satisfying $X \in [a, b]$ and $Y \in [c, d]$ almost surely, where $0 < a \leq b$ and $0 < c \leq d$. Then*

$$1 - \frac{\mathbb{E}[X] \mathbb{E}[Y]}{\mathbb{E}[XY]} \leq \frac{(b-a)(d-c)}{(\sqrt{ac} + \sqrt{bd})^2}. \quad (\text{EC.30})$$

Proof of Lemma EC.6. If $a = b$ or $c = d$, then X or Y is almost surely constant and both sides of (EC.30) vanish; we therefore assume $a < b$ and $c < d$. Let $x := \mathbb{E}[X]$ and $y := \mathbb{E}[Y]$. Since $(b - X)(Y - c) \geq 0$ and $(X - a)(d - Y) \geq 0$, taking expectations and dividing by xy gives

$$\frac{\mathbb{E}[XY]}{xy} \leq B_1(x, y) := \frac{by + cx - bc}{xy}, \quad \frac{\mathbb{E}[XY]}{xy} \leq B_2(x, y) := \frac{dx + ay - ad}{xy}.$$

Thus $\mathbb{E}[XY]/xy \leq \min\{B_1(x, y), B_2(x, y)\}$. Write $u := (x - a)/(b - a)$ and $v := (y - c)/(d - c)$. A direct comparison shows that $B_1(x, y) \leq B_2(x, y)$ if $v \leq u$, and $B_2(x, y) \leq B_1(x, y)$ if $u \leq v$. Moreover, B_1 is nondecreasing in y for fixed x , since $\partial_y B_1 = c(b - x)/(xy^2) \geq 0$, and B_2 is nondecreasing in x for fixed y , since $\partial_x B_2 = a(d - y)/(x^2 y) \geq 0$. Hence the largest possible value of $\min\{B_1, B_2\}$ occurs on the diagonal $u = v = t$.

On this diagonal, $x = a + (b - a)t$ and $y = c + (d - c)t$, and the two bounds coincide:

$$\frac{\mathbb{E}[XY]}{\mathbb{E}[X]\mathbb{E}[Y]} \leq g(t) := \frac{ac + (bd - ac)t}{(a + (b - a)t)(c + (d - c)t)}.$$

Differentiating gives

$$g'(t) = \frac{(b - a)(d - c)(ac(1 - t)^2 - bd t^2)}{(a + (b - a)t)^2 (c + (d - c)t)^2}.$$

Thus g is maximized at $t^* = \sqrt{ac}/(\sqrt{ac} + \sqrt{bd})$, where

$$g(t^*) = \frac{(\sqrt{ac} + \sqrt{bd})^2}{(\sqrt{ad} + \sqrt{bc})^2}.$$

Since $\mathbb{E}[XY]/(\mathbb{E}[X]\mathbb{E}[Y]) \leq g(t^*)$, we have $\mathbb{E}[X]\mathbb{E}[Y]/\mathbb{E}[XY] \geq 1/g(t^*)$. Therefore

$$1 - \frac{\mathbb{E}[X]\mathbb{E}[Y]}{\mathbb{E}[XY]} \leq 1 - \frac{1}{g(t^*)} = \frac{(\sqrt{ac} + \sqrt{bd})^2 - (\sqrt{ad} + \sqrt{bc})^2}{(\sqrt{ac} + \sqrt{bd})^2} = \frac{(b - a)(d - c)}{(\sqrt{ac} + \sqrt{bd})^2}.$$

□

Proof of Theorem 3. We use the notation introduced before Lemma EC.5. The proof proceeds in two steps: the first step establishes a spectral bound on the multiplicative suboptimality using Lemmas EC.2 and EC.5; and the second step maximizes this spectral bound over admissible spectra using Lemma EC.6, which yields the bound (6).

For the first step, recall that the multiplicative suboptimality is the ratio of the worst-case additive suboptimality of the shrinkage path heuristic to that of the SAA-RO heuristic. Since ρ_{gap} maximizes $\Delta(\rho)$ over $\rho \geq 0$, this ratio is

$$\frac{\sup_{\rho \geq 0} \Delta(\rho)}{\sup_{\rho \geq 0} \Delta^{0, \infty}(\rho)} = \frac{\Delta(\rho_{\text{gap}})}{\sup_{\rho \geq 0} \Delta^{0, \infty}(\rho)}, \quad (\text{EC.31})$$

where the denominator is shown below to be positive whenever $\mathbf{x}(0) \neq \mathbf{0}$.

If $\mathbf{x}(0) = \mathbf{0}$, then the WDRO solution, the shrinkage path, and the SAA-RO heuristic are all identically $\mathbf{0}$, so both additive gaps are zero for every $\rho \geq 0$; we interpret the multiplicative suboptimality as zero in this case, and the bound holds trivially. We therefore assume $\mathbf{x}(0) \neq \mathbf{0}$ in the remainder.

Lemma EC.1 gives $\mathbf{x}(\infty) = \mathbf{0}$ and the quadratic representation of f_ρ . Evaluating it at the two endpoints gives the SAA-RO value

$$\min\{f_\rho(\mathbf{x}(0)), f_\rho(\mathbf{x}(\infty))\} = f_0^* + \min\{\rho\|\mathbf{x}(0)\|_2, \mathbf{x}(0)^\top \mathbf{H}\mathbf{x}(0)\}. \quad (\text{EC.32})$$

In particular, for $\rho_{\text{end}} = \mathbf{x}(0)^\top \mathbf{H}\mathbf{x}(0)/\|\mathbf{x}(0)\|_2$, the endpoint heuristic value is $f_0^* + \mathbf{x}(0)^\top \mathbf{H}\mathbf{x}(0)$, while (EC.13) gives $\hat{f}_{\rho_{\text{end}}} = f_0^* + \frac{3}{4}\mathbf{x}(0)^\top \mathbf{H}\mathbf{x}(0)$. Hence the denominator in (EC.31) is positive:

$$\begin{aligned} \sup_{\rho \geq 0} \Delta^{0,\infty}(\rho) &\geq \Delta^{0,\infty}(\rho_{\text{end}}) = \min\{f_{\rho_{\text{end}}}(\mathbf{x}(0)), f_{\rho_{\text{end}}}(\mathbf{x}(\infty))\} - f_{\rho_{\text{end}}}^* \\ &\geq \min\{f_{\rho_{\text{end}}}(\mathbf{x}(0)), f_{\rho_{\text{end}}}(\mathbf{x}(\infty))\} - \hat{f}_{\rho_{\text{end}}} = \frac{1}{4}\mathbf{x}(0)^\top \mathbf{H}\mathbf{x}(0) > 0, \end{aligned}$$

where the last inequality uses $\mathbf{H} \succ \mathbf{0}$ and $\mathbf{x}(0) \neq \mathbf{0}$.

If $\Delta(\rho_{\text{gap}}) = 0$, then the bound is immediate. We therefore assume $\Delta(\rho_{\text{gap}}) > 0$. Lemma EC.5 gives $\rho_{\text{gap}} > \rho_{\text{end}}$, so (EC.32) gives $\min\{f_{\rho_{\text{gap}}}(\mathbf{x}(0)), f_{\rho_{\text{gap}}}(\mathbf{x}(\infty))\} = f_0^* + \mathbf{x}(0)^\top \mathbf{H}\mathbf{x}(0)$. Combining this with the identity $f_{\rho_{\text{gap}}}^* = f_0^* + \mathbf{x}(0)^\top \mathbf{H}\mathbf{x}(0) - \mathbf{x}^*(\rho_{\text{gap}})^\top \mathbf{H}\mathbf{x}^*(\rho_{\text{gap}})$ from Lemma EC.2 lower bounds the denominator of (EC.31):

$$\sup_{\rho \geq 0} \Delta^{0,\infty}(\rho) \geq \Delta^{0,\infty}(\rho_{\text{gap}}) = \mathbf{x}^*(\rho_{\text{gap}})^\top \mathbf{H}\mathbf{x}^*(\rho_{\text{gap}}).$$

Combining this lower bound with (EC.18) gives

$$\frac{\sup_{\rho \geq 0} \Delta(\rho)}{\sup_{\rho \geq 0} \Delta^{0,\infty}(\rho)} \leq \frac{\Delta(\rho_{\text{gap}})}{\Delta^{0,\infty}(\rho_{\text{gap}})} = 1 - \frac{\|\mathbf{x}^*(\rho_{\text{gap}})\|_2^2}{\|\mathbf{x}(0)\|_2^2} \frac{\mathbf{x}(0)^\top \mathbf{H}\mathbf{x}(0)}{\mathbf{x}^*(\rho_{\text{gap}})^\top \mathbf{H}\mathbf{x}^*(\rho_{\text{gap}})}.$$

Using the spectral notation of Lemma EC.2 and $\tilde{x}_i^*(\rho_{\text{gap}}) = \lambda_i(\lambda_i + \gamma_{\text{gap}})^{-1}\tilde{x}_i(0)$, the last bound becomes

$$\frac{\sup_{\rho \geq 0} \Delta(\rho)}{\sup_{\rho \geq 0} \Delta^{0,\infty}(\rho)} \leq 1 - \frac{\mathbb{E}_p[\Lambda] \mathbb{E}_p[\phi_{\gamma_{\text{gap}}}(\Lambda)]}{\mathbb{E}_p[\Lambda \phi_{\gamma_{\text{gap}}}(\Lambda)]}. \quad (\text{EC.33})$$

For the second step, we maximize the spectral representation in (EC.33). Since $\phi_{\gamma_{\text{gap}}}$ is increasing, Lemma EC.6 applies with $X = \Lambda$ and $Y = \phi_{\gamma_{\text{gap}}}(\Lambda)$. Thus, with $a = \lambda_{\min}$, $b = \lambda_{\max}$, $c = \phi_{\gamma_{\text{gap}}}(\lambda_{\min})$, and $d = \phi_{\gamma_{\text{gap}}}(\lambda_{\max})$, we obtain

$$\frac{\sup_{\rho \geq 0} \Delta(\rho)}{\sup_{\rho \geq 0} \Delta^{0,\infty}(\rho)} \leq \frac{(\lambda_{\max} - \lambda_{\min})[\phi_{\gamma_{\text{gap}}}(\lambda_{\max}) - \phi_{\gamma_{\text{gap}}}(\lambda_{\min})]}{(\lambda_{\min}^{3/2}/(\lambda_{\min} + \gamma_{\text{gap}}) + \lambda_{\max}^{3/2}/(\lambda_{\max} + \gamma_{\text{gap}}))^2} = \frac{\eta(\kappa - 1)^2(2\kappa + \eta(\kappa + 1))}{(\kappa + \eta + \kappa^{3/2}(1 + \eta))^2}, \quad (\text{EC.34})$$

where $\eta := \gamma_{\text{gap}}/\lambda_{\min}$. It remains to maximize the η -parameterized bound over $\eta \geq 0$. For fixed $\kappa \geq 1$, write this bound as $R_\kappa(\eta) := \frac{\eta(\kappa - 1)^2(2\kappa + \eta(\kappa + 1))}{(\kappa + \eta + \kappa^{3/2}(1 + \eta))^2}$. Differentiating $R_\kappa(\eta)$ gives

$$\frac{dR_\kappa(\eta)}{d\eta} = \frac{(\kappa - 1)^2(2\kappa(\kappa + \kappa^{3/2}) + 2\kappa(\kappa + \sqrt{\kappa})\eta)}{(\kappa + \eta + \kappa^{3/2}(1 + \eta))^3} \geq 0 \quad \forall \kappa \geq 1, \eta \geq 0.$$

Thus $R_\kappa(\eta)$ is nondecreasing on $[0, \infty)$, and hence

$$\sup_{\eta \geq 0} R_\kappa(\eta) = \lim_{\eta \rightarrow \infty} R_\kappa(\eta) = \frac{(\kappa - 1)^2(\kappa + 1)}{(\kappa^{3/2} + 1)^2}.$$

Combining this supremum with (EC.34) yields the bound (6) in Theorem 3. Finally, since the denominator of (6) is larger than the numerator by

$$(\kappa^{3/2} + 1)^2 - (\kappa - 1)^2(\kappa + 1) = \kappa(\sqrt{\kappa} + 1)^2 > 0,$$

the bound (6) is strictly less than one for every $\kappa \geq 1$. $\square$

Proof of Theorem 4. Throughout, write $D := W(\hat{\mathbb{P}}, \mathbb{Q}(\infty))$ and $\mathbb{Q}_\lambda := (1 - \lambda)\hat{\mathbb{P}} + \lambda\mathbb{Q}(\infty)$, so that $\widehat{\mathcal{B}}(\hat{\mathbb{P}}) = \{\mathbb{Q}_\lambda : \lambda \in [0, 1]\}$. The existence and finite support of $\mathbb{Q}(\infty)$ are established in the GitHub companion. The proof proceeds in three steps: the first step handles the trivial case $\hat{\mathbb{P}} = \mathbb{Q}(\infty)$ and otherwise shows that (9) is a maximization over those $\lambda \in [0, 1]$ for which $\mathbb{Q}_\lambda$ is feasible; the second step shows that the objective of this maximization is nondecreasing in λ , so that the largest feasible λ is optimal and yields $\widehat{\mathbb{Q}}(\rho)$; and the third step verifies that the infimum in (10) is attained.

For the first step, we begin with the trivial case $\hat{\mathbb{P}} = \mathbb{Q}(\infty)$. In this case, $\mathbb{Q}_\lambda = \hat{\mathbb{P}}$ for every $\lambda \in [0, 1]$, so the dual shrinkage path collapses to $\{\hat{\mathbb{P}}\}$ and $\widehat{\mathbb{Q}}(\rho) = \hat{\mathbb{P}}$ for every $\rho \geq 0$. Since $W(\hat{\mathbb{P}}, \hat{\mathbb{P}}) = 0 \leq \rho$, the intersection in (9) equals $\{\hat{\mathbb{P}}\}$, and thus $\underline{J}(\rho) = \min_{\mathbf{x} \in \mathcal{X}(\Xi)} \mathbb{E}_{\hat{\mathbb{P}}}[\ell(\mathbf{x}, \tilde{\boldsymbol{\xi}})]$, attained at the SAA solution $\mathbf{x}(0)$. We therefore assume $\hat{\mathbb{P}} \neq \mathbb{Q}(\infty)$ in the remainder of the proof, so that $D > 0$ since the Wasserstein distance is a metric.

Kantorovich–Rubinstein duality for the type-1 Wasserstein distance gives, for every $\lambda \in [0, 1]$,

$$W(\hat{\mathbb{P}}, \mathbb{Q}_\lambda) = \sup_{\|f\|_{\text{Lip}} \leq 1} \left\{ \int f d\hat{\mathbb{P}} - \int f d\mathbb{Q}_\lambda \right\} = \lambda \sup_{\|f\|_{\text{Lip}} \leq 1} \left\{ \int f d\hat{\mathbb{P}} - \int f d\mathbb{Q}(\infty) \right\} = \lambda D.$$

Since $\mathcal{P}(\Xi)$ is convex and $\hat{\mathbb{P}}, \mathbb{Q}(\infty) \in \mathcal{P}(\Xi)$, every $\mathbb{Q}_\lambda$ belongs to $\mathcal{P}(\Xi)$. Hence $\mathbb{Q}_\lambda \in \mathbb{B}_\rho(\hat{\mathbb{P}})$ if and only if $W(\hat{\mathbb{P}}, \mathbb{Q}_\lambda) \leq \rho$, which is equivalent to $\lambda D \leq \rho$, that is, to $\lambda \leq \rho/D$. Intersecting with $\lambda \in [0, 1]$ yields $\widehat{\mathcal{B}}(\hat{\mathbb{P}}) \cap \mathbb{B}_\rho(\hat{\mathbb{P}}) = \{\mathbb{Q}_\lambda : \lambda \in [0, \bar{\lambda}]\}$ with $\bar{\lambda} := \min\{\rho/D, 1\}$. Therefore, the a posteriori lower bound (9) is equivalent to the univariate maximization problem $\underline{J}(\rho) = \max_{\lambda \in [0, \bar{\lambda}]} \phi(\lambda)$, where $\phi(\lambda) := \inf_{\mathbf{x} \in \mathcal{X}(\Xi)} \mathbb{E}_{\mathbb{Q}_\lambda}[\ell(\mathbf{x}, \tilde{\boldsymbol{\xi}})]$.

For the second step, fix $\mathbf{x}$. The map $\lambda \mapsto \mathbb{E}_{\mathbb{Q}_\lambda}[\ell(\mathbf{x}, \tilde{\boldsymbol{\xi}})]$ is affine, so ϕ is concave on $[0, 1]$ as a pointwise infimum of affine functions. Also, $\mathbb{Q}_1 = \mathbb{Q}(\infty)$ maximizes $\inf_{\mathbf{x} \in \mathcal{X}(\Xi)} \mathbb{E}_{\mathbb{Q}}[\ell(\mathbf{x}, \tilde{\boldsymbol{\xi}})]$ over $\mathcal{P}(\Xi)$ by (8); hence $\phi(\lambda) \leq \phi(1)$ for every $\lambda \in [0, 1]$. If $0 \leq \lambda_1 < \lambda_2 \leq 1$ and $t := (\lambda_2 - \lambda_1)/(1 - \lambda_1)$, then $\lambda_2 = (1 - t)\lambda_1 + t$ and concavity gives $\phi(\lambda_2) \geq (1 - t)\phi(\lambda_1) + t\phi(1) \geq \phi(\lambda_1)$. Thus ϕ is nondecreasing, and the maximum is attained at $\bar{\lambda} = \min\{\rho/D, 1\} = 1 - \mu(\rho)$. Substituting this value into $\mathbb{Q}_{\bar{\lambda}}$ gives $\mathbb{Q}_{\bar{\lambda}} = \mu(\rho)\hat{\mathbb{P}} + (1 - \mu(\rho))\mathbb{Q}(\infty) = \widehat{\mathbb{Q}}(\rho)$, and therefore $\underline{J}(\rho) = \inf_{\mathbf{x} \in \mathcal{X}(\Xi)} \mathbb{E}_{\widehat{\mathbb{Q}}(\rho)}[\ell(\mathbf{x}, \tilde{\boldsymbol{\xi}})]$.

For the third step, it remains to replace the infimum in this representation of $\underline{J}(\rho)$ by a minimum. Since $\mathbb{Q}(\infty)$ is an RO distribution, an auxiliary result in the GitHub companion gives $\inf_{\mathbf{x} \in \mathcal{X}(\Xi)} \mathbb{E}_{\mathbb{Q}(\infty)}[\ell(\mathbf{x}, \tilde{\boldsymbol{\xi}})] = J(\infty)$, where $J(\infty)$ is the finite optimal value of the RO problem (3b). Thus $\mathbb{E}_{\mathbb{Q}(\infty)}[\ell(\mathbf{x}, \tilde{\boldsymbol{\xi}})] \geq J(\infty)$ for every $\mathbf{x} \in \mathcal{X}(\Xi)$. If $\bar{\lambda} = 1$, then $\hat{\mathbb{Q}}(\rho) = \mathbb{Q}(\infty)$, and the RO solution $\mathbf{x}(\infty)$ satisfies $\mathbb{E}_{\mathbb{Q}(\infty)}[\ell(\mathbf{x}(\infty), \tilde{\boldsymbol{\xi}})] \leq \sup_{\boldsymbol{\xi} \in \Xi} \ell(\mathbf{x}(\infty), \boldsymbol{\xi}) = J(\infty)$, so $\mathbf{x}(\infty)$ attains the infimum. If $\bar{\lambda} < 1$, then $\mu(\rho) > 0$. For every $\mathbf{x} \in \mathcal{X}(\Xi)$,

$$\begin{aligned} \mathbb{E}_{\hat{\mathbb{Q}}(\rho)}[\ell(\mathbf{x}, \tilde{\boldsymbol{\xi}})] &= \mu(\rho) \mathbb{E}_{\hat{\mathbb{P}}}[\ell(\mathbf{x}, \tilde{\boldsymbol{\xi}})] + (1 - \mu(\rho)) \mathbb{E}_{\mathbb{Q}(\infty)}[\ell(\mathbf{x}, \tilde{\boldsymbol{\xi}})] \\ &\geq \mu(\rho) \mathbb{E}_{\hat{\mathbb{P}}}[\ell(\mathbf{x}, \tilde{\boldsymbol{\xi}})] + (1 - \mu(\rho)) J(\infty). \end{aligned}$$

Assumption 3(i) makes the right-hand side coercive on $\mathcal{X}(\Xi)$, and hence the objective is coercive as well. The distribution $\hat{\mathbb{Q}}(\rho)$ is finitely supported, so the objective is a finite nonnegative weighted sum of lower semicontinuous functions $\ell(\cdot, \boldsymbol{\xi})$ and is therefore lower semicontinuous. The set $\mathcal{X}(\Xi)$ is closed by the closedness of $\mathcal{X}$ and of each constraint $g_m(\cdot, \boldsymbol{\xi}) \leq \gamma_m$. Weierstrass' theorem therefore gives an optimizer, and the infimum is a minimum. $\square$

Proof of Theorem 5. We prove the three claims in order. Throughout, write $D := W(\hat{\mathbb{P}}, \mathbb{Q}(\infty))$.

For monotonicity, let $0 \leq \rho_1 \leq \rho_2$. Then $\mathbb{B}_{\rho_1}(\hat{\mathbb{P}}) \subseteq \mathbb{B}_{\rho_2}(\hat{\mathbb{P}})$. For every fixed $\mathbf{x}$, the inner worst-case expectation $\sup_{\mathbb{Q} \in \mathbb{B}_{\rho}(\hat{\mathbb{P}})} \mathbb{E}_{\mathbb{Q}}[\ell(\mathbf{x}, \tilde{\boldsymbol{\xi}})]$ is therefore nondecreasing in ρ . Taking the minimum over the ρ -independent sets $\mathcal{X}(\Xi)$ and $\hat{\mathcal{X}}(\Xi)$ shows that $J^*(\rho)$ and $\bar{J}(\rho)$ are nondecreasing. For the lower bound, the feasible sets in (9) also nest: $\hat{\mathcal{B}}(\hat{\mathbb{P}}) \cap \mathbb{B}_{\rho_1}(\hat{\mathbb{P}}) \subseteq \hat{\mathcal{B}}(\hat{\mathbb{P}}) \cap \mathbb{B}_{\rho_2}(\hat{\mathbb{P}})$. Maximizing the same objective over a larger set cannot decrease the value, so $\underline{J}(\rho)$ is nondecreasing.

For concavity, we first consider J^* and $\bar{J}$. The strong duality of the Wasserstein worst-case expectation (Mohajerin Esfahani and Kuhn 2018, Theorem 4.2) gives, for every fixed $\mathbf{x}$ and every radius $\rho > 0$,

$$\sup_{\mathbb{Q} \in \mathbb{B}_{\rho}(\hat{\mathbb{P}})} \mathbb{E}_{\mathbb{Q}}[\ell(\mathbf{x}, \tilde{\boldsymbol{\xi}})] = \inf_{\eta \geq 0} \left\{ \eta \rho + \frac{1}{N} \sum_{i=1}^N \sup_{\boldsymbol{\xi} \in \Xi} [\ell(\mathbf{x}, \boldsymbol{\xi}) - \eta d(\boldsymbol{\xi}, \hat{\boldsymbol{\xi}}_i)] \right\}.$$

Thus the worst-case expectation is an infimum of functions affine in ρ . Taking the further infimum over $\mathbf{x} \in \mathcal{X}(\Xi)$ gives $J^*(\rho)$, and taking it over $\mathbf{x} \in \hat{\mathcal{X}}(\Xi)$ gives $\bar{J}(\rho)$. Since both feasible sets are independent of ρ , both value functions are pointwise infima of affine functions of ρ , and hence concave on $(0, \infty)$. Concavity on $[0, \infty)$ follows because both functions are nondecreasing by claim (i), and a nondecreasing function on $[0, \infty)$ that is concave on $(0, \infty)$ is concave throughout. For $\underline{J}$, if $D = 0$, then $\underline{J}$ is constant—and hence concave—by Theorem 4. If $D > 0$, the proof of Theorem 4 shows that $\underline{J}(\rho) = \phi(\min\{\rho/D, 1\})$, where $\phi(\lambda) := \inf_{\mathbf{x} \in \mathcal{X}(\Xi)} \mathbb{E}_{(1-\lambda)\hat{\mathbb{P}} + \lambda\mathbb{Q}(\infty)}[\ell(\mathbf{x}, \tilde{\boldsymbol{\xi}})]$ is concave and nondecreasing on $[0, 1]$. The map $\rho \mapsto \min\{\rho/D, 1\}$ is concave, and composing a nondecreasing concave function with a concave function preserves concavity. Hence $\underline{J}$ is concave.

For endpoint tightness, first take $\rho = 0$. The Wasserstein ball is then the singleton $\mathbb{B}_0(\hat{\mathbb{P}}) = \{\hat{\mathbb{P}}\}$, so $J^*(0) = \min_{\mathbf{x} \in \mathcal{X}(\Xi)} \mathbb{E}_{\hat{\mathbb{P}}}[\ell(\mathbf{x}, \tilde{\boldsymbol{\xi}})] = J(0)$, which equals the optimal value of the SAA problem (3a).

Since $\mathbf{x}(0) \in \hat{\mathcal{X}}(\Xi)$ is SAA-optimal over the larger set $\mathcal{X}(\Xi)$, we also have $\bar{J}(0) = J(0)$. The dual shrinkage path intersects $\mathbb{B}_0(\hat{\mathbb{P}})$ only at $\hat{\mathbb{P}}$, so $\underline{J}(0) = J(0)$ as well. Thus $\bar{J}(0) = J^*(0) = \underline{J}(0) = J(0)$.

Finally, let $\rho \geq D$. Theorem 4 gives $\hat{\mathbb{Q}}(\rho) = \mathbb{Q}(\infty)$, with the degenerate case $D = 0$ covered by $\hat{\mathbb{P}} = \mathbb{Q}(\infty)$, so $\underline{J}(\rho) = \min_{\mathbf{x} \in \mathcal{X}(\Xi)} \mathbb{E}_{\mathbb{Q}(\infty)}[\ell(\mathbf{x}, \tilde{\boldsymbol{\xi}})]$, which equals $J(\infty)$ by the same auxiliary result from the GitHub companion invoked in the proof of Theorem 4. Note that $\mathbb{B}_\rho(\hat{\mathbb{P}}) \subseteq \mathcal{P}(\Xi)$ implies

$$J^*(\rho) \leq \min_{\mathbf{x} \in \mathcal{X}(\Xi)} \sup_{\boldsymbol{\xi} \in \Xi} \ell(\mathbf{x}, \boldsymbol{\xi}) = J(\infty).$$

Since also $\underline{J}(\rho) \leq J^*(\rho)$ by (9) and $\underline{J}(\rho) = J(\infty)$, we have $J^*(\rho) = J(\infty)$. For the upper bound, evaluating the path problem at $\mathbf{x}(\infty) \in \hat{\mathcal{X}}(\Xi)$ gives $\bar{J}(\rho) \leq \sup_{\mathbb{Q} \in \mathbb{B}_\rho(\hat{\mathbb{P}})} \mathbb{E}_{\mathbb{Q}}[\ell(\mathbf{x}(\infty), \tilde{\boldsymbol{\xi}})] \leq J(\infty)$. The opposite inequality $\bar{J}(\rho) \geq J^*(\rho) = J(\infty)$ follows from $\hat{\mathcal{X}}(\Xi) \subseteq \mathcal{X}(\Xi)$. Hence $\bar{J}(\rho) = J^*(\rho) = \underline{J}(\rho) = J(\infty)$ for every $\rho \geq D$. $\square$

The proof of Proposition 2 relies on the following three lemmas, one for each of the settings (i)–(iii). Each lemma exhibits a finite convex problem whose solution yields a finitely supported RO distribution $\mathbb{Q}(\infty)$ defined in (8).

LEMMA EC.7 (Vertex Construction). *Let the support $\Xi = \text{conv}(\mathbf{v}_1, \dots, \mathbf{v}_K) + \text{cone}(\mathbf{r}_1, \dots, \mathbf{r}_R)$ be a pointed polyhedron with vertices $\mathbf{v}_1, \dots, \mathbf{v}_K$ and extreme rays $\mathbf{r}_1, \dots, \mathbf{r}_R$, and let the loss $\ell(\mathbf{x}, \cdot)$ be convex on Ξ with nonpositive recession along every extreme ray, that is,*

$$\ell^\infty(\mathbf{x}, \mathbf{r}_r) := \lim_{\alpha \rightarrow \infty} \alpha^{-1} [\ell(\mathbf{x}, \boldsymbol{\xi} + \alpha \mathbf{r}_r) - \ell(\mathbf{x}, \boldsymbol{\xi})] \leq 0 \quad \forall \mathbf{x} \in \mathcal{X}(\Xi), \boldsymbol{\xi} \in \Xi, r \in [R].$$

Then the RO problem (3b) admits the finite vertex reformulation

$$J(\infty) = \min_{\mathbf{x} \in \mathcal{X}(\Xi), t} \{t : \ell(\mathbf{x}, \mathbf{v}_j) \leq t, j \in [K]\}, \quad (\text{EC.35})$$

and a finitely supported RO distribution can be obtained from its dual.

Proof of Lemma EC.7. Fix $\mathbf{x} \in \mathcal{X}(\Xi)$ and write any $\boldsymbol{\xi} \in \Xi$ as $\boldsymbol{\xi} = \bar{\boldsymbol{\xi}} + \sum_r \beta_r \mathbf{r}_r$ with $\bar{\boldsymbol{\xi}} = \sum_j \alpha_j \mathbf{v}_j$, $\boldsymbol{\alpha} \in \Delta_K$, and $\beta_r \geq 0$, where Δ_K denotes the probability simplex on $[K]$. Nonpositive recession gives $\ell(\mathbf{x}, \boldsymbol{\xi}) \leq \ell(\mathbf{x}, \bar{\boldsymbol{\xi}})$, while convexity gives $\ell(\mathbf{x}, \bar{\boldsymbol{\xi}}) \leq \sum_j \alpha_j \ell(\mathbf{x}, \mathbf{v}_j) \leq \max_{j \in [K]} \ell(\mathbf{x}, \mathbf{v}_j)$; since each vertex lies in Ξ , the reverse inequality is immediate, so $\sup_{\boldsymbol{\xi} \in \Xi} \ell(\mathbf{x}, \boldsymbol{\xi}) = \max_{j \in [K]} \ell(\mathbf{x}, \mathbf{v}_j)$. The RO problem (3b) therefore reduces to the finite vertex reformulation (EC.35).

Dualizing the epigraph constraints of (EC.35) with multipliers $\boldsymbol{\pi} \geq \mathbf{0}$ and minimizing the Lagrangian over t forces $\sum_j \pi_j = 1$, giving the dual $\max_{\boldsymbol{\pi} \in \Delta_K} \inf_{\mathbf{x} \in \mathcal{X}(\Xi)} \sum_j \pi_j \ell(\mathbf{x}, \mathbf{v}_j)$. As any $t > J(\infty)$ is strictly feasible at $\mathbf{x}(\infty)$, strong duality holds and an optimal $\boldsymbol{\pi}^* \in \Delta_K$ exists with $\inf_{\mathbf{x} \in \mathcal{X}(\Xi)} \sum_j \pi_j^* \ell(\mathbf{x}, \mathbf{v}_j) = J(\infty)$. Placing the optimal mass on the active vertices yields the candidate RO distribution

$$\mathbb{Q}(\infty) = \sum_{j: \pi_j^* > 0} \pi_j^* \delta_{\mathbf{v}_j}.$$

This is indeed an RO distribution: $\inf_{\mathbf{x} \in \mathcal{X}(\Xi)} \mathbb{E}_{\mathbb{Q}(\infty)}[\ell(\mathbf{x}, \tilde{\boldsymbol{\xi}})] = \inf_{\mathbf{x} \in \mathcal{X}(\Xi)} \sum_j \pi_j^* \ell(\mathbf{x}, \mathbf{v}_j) = J(\infty)$, whereas every $\mathbb{Q} \in \mathcal{P}(\Xi)$ satisfies $\inf_{\mathbf{x} \in \mathcal{X}(\Xi)} \mathbb{E}_{\mathbb{Q}}[\ell(\mathbf{x}, \tilde{\boldsymbol{\xi}})] \leq \mathbb{E}_{\mathbb{Q}}[\ell(\mathbf{x}(\infty), \tilde{\boldsymbol{\xi}})] \leq \sup_{\boldsymbol{\xi} \in \Xi} \ell(\mathbf{x}(\infty), \boldsymbol{\xi}) = J(\infty)$, so $\mathbb{Q}(\infty)$ attains the maximum in (8) and is finitely supported. $\square$

LEMMA EC.8 (Finite-Max Construction). *In addition to the assumptions of Lemma EC.7, let the feasible set be the polyhedron $\mathcal{X}(\Xi) = \{\mathbf{x} \in \mathbb{R}^n : C\mathbf{x} \leq \mathbf{d}\}$, and let $\ell(\mathbf{x}, \boldsymbol{\xi}) = \max_{l \in [L]} \ell_l(\mathbf{x}, \boldsymbol{\xi})$, where each $\ell_l(\cdot, \boldsymbol{\xi})$ is convex and differentiable. Define the active vertices, the active loss pieces at each vertex, and the active constraints at the RO solution $(\mathbf{x}(\infty), J(\infty))$ as*

$$\mathcal{A} = \{j : \ell(\mathbf{x}(\infty), \mathbf{v}_j) = J(\infty)\}, \quad \mathcal{L}_j = \{l : \ell_l(\mathbf{x}(\infty), \mathbf{v}_j) = J(\infty)\}, \quad \mathcal{I}_\infty = \{m : C_m \mathbf{x}(\infty) = d_m\}.$$

Then the feasibility linear program

$$\sum_{j \in \mathcal{A}} \sum_{l \in \mathcal{L}_j} \pi_{jl} \nabla_{\mathbf{x}} \ell_l(\mathbf{x}(\infty), \mathbf{v}_j) + C_{\mathcal{I}_\infty}^\top \boldsymbol{\mu} = \mathbf{0}, \quad \sum_{j \in \mathcal{A}} \sum_{l \in \mathcal{L}_j} \pi_{jl} = 1, \quad \pi_{jl} \geq 0, \quad \boldsymbol{\mu} \geq \mathbf{0} \quad (\text{EC.36})$$

is feasible, and every basic feasible solution induces an RO distribution supported on at most $n+1$ atoms.

Proof of Lemma EC.8. The argument in Lemma EC.7 gives $\sup_{\boldsymbol{\xi} \in \Xi} \ell(\mathbf{x}, \boldsymbol{\xi}) = \max_{j \in [K]} \max_{l \in [L]} \ell_l(\mathbf{x}, \mathbf{v}_j)$, so $\mathbf{x}(\infty)$ minimizes this finite maximum over $\mathcal{X}(\Xi)$, and the active terms of the maximum at $\mathbf{x}(\infty)$ are exactly the pairs $\{(j, l) : j \in \mathcal{A}, l \in \mathcal{L}_j\}$.

The linear program (EC.36) is feasible. Indeed, since $\mathbf{x}(\infty)$ minimizes the finite maximum over the polyhedron $\mathcal{X}(\Xi)$, first-order optimality gives $\mathbf{0} \in \partial_{\mathbf{x}}(\max_{j,l} \ell_l(\cdot, \mathbf{v}_j))(\mathbf{x}(\infty)) + N_{\mathcal{X}(\Xi)}(\mathbf{x}(\infty))$, where $N_{\mathcal{X}(\Xi)}(\mathbf{x}) := \{\mathbf{z} \in \mathbb{R}^n : \mathbf{z}^\top(\mathbf{y} - \mathbf{x}) \leq 0 \ \forall \mathbf{y} \in \mathcal{X}(\Xi)\}$ denotes the normal cone to $\mathcal{X}(\Xi)$ at $\mathbf{x}$. The subdifferential of a finite maximum of differentiable convex functions is the convex hull of the active gradients, $\text{conv}\{\nabla_{\mathbf{x}} \ell_l(\mathbf{x}(\infty), \mathbf{v}_j) : j \in \mathcal{A}, l \in \mathcal{L}_j\}$, and the normal cone to the polyhedron is $\{C_{\mathcal{I}_\infty}^\top \boldsymbol{\mu} : \boldsymbol{\mu} \geq \mathbf{0}\}$; combining the two identities produces nonnegative π_{jl} summing to one together with a $\boldsymbol{\mu} \geq \mathbf{0}$ that satisfy (EC.36).

Any feasible solution of (EC.36) induces an RO distribution. Let $(\pi_{jl}^S, \boldsymbol{\mu}^S)$ be such a solution, put $\pi_j^S := \sum_{l \in \mathcal{L}_j} \pi_{jl}^S$, and define $\mathbb{Q}^S(\infty) := \sum_{j \in \mathcal{A}: \pi_j^S > 0} \pi_j^S \delta_{\mathbf{v}_j}$. For each such j the averaged gradient $\mathbf{g}_j^S := (\pi_j^S)^{-1} \sum_{l \in \mathcal{L}_j} \pi_{jl}^S \nabla_{\mathbf{x}} \ell_l(\mathbf{x}(\infty), \mathbf{v}_j)$ is a convex combination of active gradients and hence lies in $\partial_{\mathbf{x}} \ell(\mathbf{x}(\infty), \mathbf{v}_j)$, while $\mathbf{z}^S := C_{\mathcal{I}_\infty}^\top \boldsymbol{\mu}^S \in N_{\mathcal{X}(\Xi)}(\mathbf{x}(\infty))$, and (EC.36) reads $\sum_{j \in \mathcal{A}} \pi_j^S \mathbf{g}_j^S + \mathbf{z}^S = \mathbf{0}$. For any $\mathbf{x} \in \mathcal{X}(\Xi)$ the subgradient inequality gives $\ell(\mathbf{x}, \mathbf{v}_j) \geq J(\infty) + (\mathbf{g}_j^S)^\top(\mathbf{x} - \mathbf{x}(\infty))$ for every $j \in \mathcal{A}$; weighting by π_j^S and summing gives

$$\mathbb{E}_{\mathbb{Q}^S(\infty)}[\ell(\mathbf{x}, \tilde{\boldsymbol{\xi}})] \geq J(\infty) + \left(\sum_{j \in \mathcal{A}} \pi_j^S \mathbf{g}_j^S \right)^\top (\mathbf{x} - \mathbf{x}(\infty)) = J(\infty) - (\mathbf{z}^S)^\top (\mathbf{x} - \mathbf{x}(\infty)) \geq J(\infty),$$

the last step because $\mathbf{z}^S \in N_{\mathcal{X}(\Xi)}(\mathbf{x}(\infty))$. Equality at $\mathbf{x} = \mathbf{x}(\infty)$ gives $\inf_{\mathbf{x} \in \mathcal{X}(\Xi)} \mathbb{E}_{\mathbb{Q}^S(\infty)}[\ell(\mathbf{x}, \tilde{\boldsymbol{\xi}})] = J(\infty)$, and since every $\mathbb{Q} \in \mathcal{P}(\Xi)$ obeys $\inf_{\mathbf{x} \in \mathcal{X}(\Xi)} \mathbb{E}_{\mathbb{Q}}[\ell(\mathbf{x}, \tilde{\boldsymbol{\xi}})] \leq \mathbb{E}_{\mathbb{Q}}[\ell(\mathbf{x}(\infty), \tilde{\boldsymbol{\xi}})] \leq \sup_{\boldsymbol{\xi} \in \Xi} \ell(\mathbf{x}(\infty), \boldsymbol{\xi}) = J(\infty)$, the distribution $\mathbb{Q}^S(\infty)$ is an RO distribution.

Finally, the bound on the number of atoms follows from the structure of a basic feasible solution. The system (EC.36) has $n+1$ equality constraints— n stationarity equations and one normalization—so a basic feasible solution has at most $n+1$ positive entries among the variables $\{\pi_{jl}\} \cup \{\boldsymbol{\mu}\}$, hence at most $n+1$ positive π_{jl}^S and therefore at most $n+1$ atoms $\mathbf{v}_j$ in $\mathbb{Q}^S(\infty)$, matching the bound established in the GitHub companion. $\square$

LEMMA EC.9 (Convex–Concave Construction). *Let the support Ξ be convex and compact, and let $\ell(\mathbf{x}, \boldsymbol{\xi}) = \max_{j \in [K]} \ell_j(\mathbf{x}, \boldsymbol{\xi})$, where each ℓ_j is closed convex with full domain in $\mathbf{x}$ and concave and upper semicontinuous in $\boldsymbol{\xi}$ on Ξ . Write $\ell_j^*(\mathbf{s}, \boldsymbol{\xi}) := \sup_{\mathbf{x} \in \mathbb{R}^n} \{\mathbf{s}^\top \mathbf{x} - \ell_j(\mathbf{x}, \boldsymbol{\xi})\}$ for the Fenchel conjugate of each piece in its first argument, $\sigma_{\mathcal{X}(\Xi)}(\mathbf{z}) := \sup_{\mathbf{x} \in \mathcal{X}(\Xi)} \mathbf{z}^\top \mathbf{x}$ for the support function of $\mathcal{X}(\Xi)$, and Δ_K for the probability simplex on $[K]$. Then the RO problem (3b) admits the finite convex reformulation*

$$J(\infty) = \max_{\boldsymbol{\pi} \in \Delta_K, \boldsymbol{\zeta}_j/\pi_j \in \Xi, \mathbf{s}_j \in \mathbb{R}^n} \left\{ -\sigma_{\mathcal{X}(\Xi)}\left(-\sum_{j=1}^K \mathbf{s}_j\right) - \sum_{j=1}^K \pi_j \ell_j^*(\mathbf{s}_j/\pi_j, \boldsymbol{\zeta}_j/\pi_j) \right\}, \quad (\text{EC.37})$$

with the standard closed-perspective convention at $\pi_j = 0$, under which the j th term of the sum equals 0 if $(\mathbf{s}_j, \boldsymbol{\zeta}_j) = (\mathbf{0}, \mathbf{0})$ and $+\infty$ otherwise, and the constraint $\boldsymbol{\zeta}_j/\pi_j \in \Xi$ reduces to $\boldsymbol{\zeta}_j = \mathbf{0}$, and every maximizer of (EC.37) induces a finitely supported RO distribution.

Proof of Lemma EC.9. For each $\mathbf{x}$ we have $\sup_{\boldsymbol{\xi} \in \Xi} \ell(\mathbf{x}, \boldsymbol{\xi}) = \max_{j \in [K]} \phi_j(\mathbf{x})$ with $\phi_j(\mathbf{x}) := \max_{\boldsymbol{\xi} \in \Xi} \ell_j(\mathbf{x}, \boldsymbol{\xi})$, where the maxima are attained because Ξ is compact and $\ell_j(\mathbf{x}, \cdot)$ is upper semicontinuous, and each ϕ_j is real-valued and convex in $\mathbf{x}$, hence continuous. Lifting the finite maximum to the simplex gives $\max_j \phi_j(\mathbf{x}) = \max_{\boldsymbol{\pi} \in \Delta_K} \sum_j \pi_j \phi_j(\mathbf{x})$, and Sion’s minimax theorem (Sion 1958) over the convex set $\mathcal{X}(\Xi)$ and the convex compact set Δ_K yields $J(\infty) = \max_{\boldsymbol{\pi} \in \Delta_K} \inf_{\mathbf{x} \in \mathcal{X}(\Xi)} \sum_j \pi_j \phi_j(\mathbf{x})$. For fixed $\boldsymbol{\pi}$, separability gives $\sum_j \pi_j \phi_j(\mathbf{x}) = \max_{\boldsymbol{\xi}_1, \dots, \boldsymbol{\xi}_K \in \Xi} \sum_j \pi_j \ell_j(\mathbf{x}, \boldsymbol{\xi}_j)$, an objective convex in $\mathbf{x}$ and concave and upper semicontinuous in $(\boldsymbol{\xi}_1, \dots, \boldsymbol{\xi}_K)$. A second application of Sion over the convex compact set Ξ^K swaps the inner infimum and maximum, and combining the two interchanges yields

$$J(\infty) = \max_{\boldsymbol{\pi} \in \Delta_K, \boldsymbol{\xi}_1, \dots, \boldsymbol{\xi}_K \in \Xi} \inf_{\mathbf{x} \in \mathcal{X}(\Xi)} \sum_{j=1}^K \pi_j \ell_j(\mathbf{x}, \boldsymbol{\xi}_j). \quad (\text{EC.38})$$

At any maximizer $(\boldsymbol{\pi}^*, \boldsymbol{\xi}_1^*, \dots, \boldsymbol{\xi}_K^*)$ the inner infimum equals $J(\infty)$ and is attained at $\mathbf{x}(\infty)$, since $\sum_j \pi_j^* \ell_j(\mathbf{x}(\infty), \boldsymbol{\xi}_j^*) \leq \sup_{\boldsymbol{\xi} \in \Xi} \ell(\mathbf{x}(\infty), \boldsymbol{\xi}) = J(\infty)$.

Fix $(\boldsymbol{\pi}, \boldsymbol{\xi}_1, \dots, \boldsymbol{\xi}_K) \in \Delta_K \times \Xi^K$ with $\boldsymbol{\pi} > \mathbf{0}$ and dualize the inner minimization of $\sum_j \pi_j \ell_j(\mathbf{x}, \boldsymbol{\xi}_j)$ over $\mathcal{X}(\Xi)$; zero-weight pieces drop from this sum and are reinstated below via the closed-perspective

convention with $(\mathbf{s}_j, \boldsymbol{\zeta}_j) = (\mathbf{0}, \mathbf{0})$. Each summand is proper closed convex with full domain $\mathbb{R}^n$, so the objective function $\sum_j \pi_j \ell_j(\cdot, \boldsymbol{\xi}_j)$ is finite convex on all of $\mathbb{R}^n$. The indicator function of $\mathcal{X}(\Xi)$ equals 0 on $\mathcal{X}(\Xi)$ and $+\infty$ outside. Adding the indicator function to the objective function enforces the constraint $\mathbf{x} \in \mathcal{X}(\Xi)$. The inner infimum in (EC.38) therefore equals the minimization of this sum over all of $\mathbb{R}^n$. Fenchel–Rockafellar duality applies to this sum of the objective function and the indicator function defined by $\mathcal{X}(\Xi)$. Its qualification requires the relative interiors of their domains, $\mathbb{R}^n$ and $\mathcal{X}(\Xi)$, to intersect. This holds since $\mathcal{X}(\Xi)$ is nonempty, and weak duality reads

$$-\sigma_{\mathcal{X}(\Xi)}\left(-\sum_j \mathbf{s}_j\right) - \sum_j \pi_j \ell_j^*(\mathbf{s}_j/\pi_j, \boldsymbol{\xi}_j) \leq \inf_{\mathbf{x} \in \mathcal{X}(\Xi)} \sum_j \pi_j \ell_j(\mathbf{x}, \boldsymbol{\xi}_j) \quad \forall \mathbf{s}_1, \dots, \mathbf{s}_K \in \mathbb{R}^n,$$

with equality and dual attainment whenever the inner infimum is finite—in particular at a maximizer of (EC.38). Because $\ell_j(\mathbf{x}, \cdot)$ is concave, $\ell_j^*(\mathbf{s}, \boldsymbol{\xi})$ is jointly convex in $(\mathbf{s}, \boldsymbol{\xi})$, so its perspective $(\pi_j, \mathbf{s}_j, \boldsymbol{\zeta}_j) \mapsto \pi_j \ell_j^*(\mathbf{s}_j/\pi_j, \boldsymbol{\zeta}_j/\pi_j)$ is jointly convex; substituting $\boldsymbol{\zeta}_j = \pi_j \boldsymbol{\xi}_j$ turns the constraint $\boldsymbol{\xi}_j \in \Xi$ into $\boldsymbol{\zeta}_j/\pi_j \in \Xi$ and, maximizing the dual jointly with (EC.38), produces the convex reformulation (EC.37). The maximization problem (EC.37) has a concave objective and a convex feasible set. Weak duality bounds its objective at every feasible point by the corresponding inner infimum in (EC.38), which is at most $J(\infty)$. At a maximizer of (EC.38), this inner infimum is finite and equals $J(\infty)$. Fenchel–Rockafellar duality therefore yields multipliers $\mathbf{s}_1, \dots, \mathbf{s}_K$ for which the objective in (EC.37) attains $J(\infty)$. Thus, the optimal value of (EC.37) is $J(\infty)$ and is attained.

Let $(\boldsymbol{\pi}^*, \boldsymbol{\zeta}_1^*, \dots, \boldsymbol{\zeta}_K^*, \mathbf{s}_1^*, \dots, \mathbf{s}_K^*)$ maximize (EC.37), and for $\pi_j^* > 0$ set $\boldsymbol{\xi}_j^* := \boldsymbol{\zeta}_j^*/\pi_j^* \in \Xi$. Fenchel–Rockafellar duality gives $\inf_{\mathbf{x} \in \mathcal{X}(\Xi)} \sum_j \pi_j^* \ell_j(\mathbf{x}, \boldsymbol{\xi}_j^*) = J(\infty)$. Define

$$\mathbb{Q}(\infty) = \sum_{j: \pi_j^* > 0} \pi_j^* \delta_{\boldsymbol{\xi}_j^*}.$$

For any $\mathbf{x} \in \mathcal{X}(\Xi)$, $\ell = \max_r \ell_r \geq \ell_j$ yields $\mathbb{E}_{\mathbb{Q}(\infty)}[\ell(\mathbf{x}, \tilde{\boldsymbol{\xi}})] = \sum_j \pi_j^* \ell(\mathbf{x}, \boldsymbol{\xi}_j^*) \geq \sum_j \pi_j^* \ell_j(\mathbf{x}, \boldsymbol{\xi}_j^*)$. Moreover, $\sum_j \pi_j^* \ell_j(\mathbf{x}, \boldsymbol{\xi}_j^*) \geq \inf_{\mathbf{y} \in \mathcal{X}(\Xi)} \sum_j \pi_j^* \ell_j(\mathbf{y}, \boldsymbol{\xi}_j^*) = J(\infty)$. Thus, $\mathbb{E}_{\mathbb{Q}(\infty)}[\ell(\mathbf{x}, \tilde{\boldsymbol{\xi}})] \geq J(\infty)$ for every $\mathbf{x} \in \mathcal{X}(\Xi)$. Taking the infimum over $\mathbf{x} \in \mathcal{X}(\Xi)$ yields $\inf_{\mathbf{x} \in \mathcal{X}(\Xi)} \mathbb{E}_{\mathbb{Q}(\infty)}[\ell(\mathbf{x}, \tilde{\boldsymbol{\xi}})] \geq J(\infty)$. Conversely, the infimum is at most $\mathbb{E}_{\mathbb{Q}(\infty)}[\ell(\mathbf{x}(\infty), \tilde{\boldsymbol{\xi}})] \leq \sup_{\boldsymbol{\xi} \in \Xi} \ell(\mathbf{x}(\infty), \boldsymbol{\xi}) = J(\infty)$, so it equals $J(\infty)$. Since every $\mathbb{Q} \in \mathcal{P}(\Xi)$ obeys $\inf_{\mathbf{x} \in \mathcal{X}(\Xi)} \mathbb{E}_{\mathbb{Q}}[\ell(\mathbf{x}, \tilde{\boldsymbol{\xi}})] \leq \mathbb{E}_{\mathbb{Q}}[\ell(\mathbf{x}(\infty), \tilde{\boldsymbol{\xi}})] \leq J(\infty)$, the distribution $\mathbb{Q}(\infty)$ attains the maximum in (8) and is a finitely supported RO distribution. $\square$

Proof of Proposition 2. Lemmas EC.7, EC.8, and EC.9 construct a finitely supported RO distribution from the solution of a finite convex problem under the assumptions of settings (i), (ii), and (iii), respectively. $\square$