

Supplementary Material to Robust Portfolio Optimization with
respect to Spectral Risk Measures under Correlation Uncertainty
(submitted to Applied Mathematics and Optimization)

Man Yiu (Tim) Tsang^{*1}, Tony Sit^{†1}, and Hoi Ying Wong^{‡1}

¹Department of Statistics, The Chinese University of Hong Kong, Shatin, N.T., Hong
Kong

A. Proofs

A.1. Proof of Convexity of $\mathcal{P}$

Lemma A.1. *The uncertainty set $\mathcal{P}$ defined in (4) is convex.*

Proof. Let π_1 and π_2 be two distributions in $\mathcal{P}$ with $\pi_\lambda = \lambda\pi_1 + (1 - \lambda)\pi_2$ for any $\lambda \in [0, 1]$. Trivially,

$$\int d\pi_\lambda = \lambda \int d\pi_1 + (1 - \lambda) \int d\pi_2 = 1 \quad \text{and} \quad \mathbb{E}_{\pi_\lambda}(\boldsymbol{\xi}) = \int \boldsymbol{\xi} d\pi_\lambda = \boldsymbol{\mu}.$$

Finally, notice that

$$\text{Var}_{\pi_\lambda}(\boldsymbol{\xi}) = \mathbb{E}_{\pi_\lambda}[(\boldsymbol{\xi} - \boldsymbol{\mu})(\boldsymbol{\xi} - \boldsymbol{\mu})^\top] = \lambda \text{Var}_{\pi_1}(\boldsymbol{\xi}) + (1 - \lambda) \text{Var}_{\pi_2}(\boldsymbol{\xi}).$$

The diagonal entries imply that $\text{Var}_{\pi_\lambda}(\xi_i) = \sigma_i^2$ for $i = 1, \dots, n$ while the off-diagonal entries guarantee the correlation bounds. The positive semidefiniteness of $\mathbf{C}_{\pi_\lambda}$ follows from the positive semidefiniteness of $\text{Var}_{\pi_\lambda}(\boldsymbol{\xi})$. Therefore, $\pi_\lambda \in \mathcal{P}$ and the set is convex. $\square$

^{*}Email: 1155063646@link.cuhk.edu.hk

[†]Corresponding author. Email: tonysit@sta.cuhk.edu.hk

[‡]Email: hywong@cuhk.edu.hk

A.2. Proof of Theorem 1

Proof of Theorem 1. For notational simplicity, we suppress the dependence of the associated distribution on the risk measure ρ_ϕ . First, define $\mathcal{R}(\boldsymbol{\mu}, \boldsymbol{\Sigma}) = \{\pi \in \mathcal{D} \mid \mathbb{E}_\pi(\boldsymbol{\xi}) = \boldsymbol{\mu}, \text{Var}_\pi(\boldsymbol{\xi}) = \boldsymbol{\sigma}^\top \mathbf{C} \boldsymbol{\sigma}, \text{Corr}_\pi(\boldsymbol{\xi}) = \mathbf{C}\}$. Note that

$$\sup_{\mathbb{P} \in \mathcal{P}} \rho_\phi(-\mathbf{a}^\top \boldsymbol{\xi}) = \sup_{\mathbf{C} \in \mathcal{C}} \sup_{\pi \in \mathcal{R}(\boldsymbol{\mu}, \boldsymbol{\sigma}^\top \mathbf{C} \boldsymbol{\sigma})} \rho_\phi(-\mathbf{a}^\top \boldsymbol{\xi})$$

(see, e.g., Kuhn et al., 2019). Then, by Lemma 2.4 of Chen et al. (2011), we have

$$\sup_{\pi \in \mathcal{R}(\boldsymbol{\mu}, \boldsymbol{\sigma}^\top \mathbf{C} \boldsymbol{\sigma})} \rho_\phi(-\mathbf{a}^\top \boldsymbol{\xi}) = \sup_{\pi \in \mathcal{R}(-\mathbf{a}^\top \boldsymbol{\mu}, \mathbf{a}^\top \boldsymbol{\sigma}^\top \mathbf{C} \boldsymbol{\sigma} \mathbf{a})} \rho_\phi(\boldsymbol{\zeta}),$$

where $\mathcal{R}(-\mathbf{a}^\top \boldsymbol{\mu}, \mathbf{a}^\top \boldsymbol{\sigma}^\top \mathbf{C} \boldsymbol{\sigma} \mathbf{a})$ is the set of univariate distribution with mean $-\mathbf{a}^\top \boldsymbol{\mu}$ and variance $\mathbf{a}^\top \boldsymbol{\sigma}^\top \mathbf{C} \boldsymbol{\sigma} \mathbf{a}$. Provoking Theorem 2 in Li (2018), we obtain

$$\sup_{\pi \in \mathcal{R}(-\mathbf{a}^\top \boldsymbol{\mu}, \mathbf{a}^\top \boldsymbol{\sigma}^\top \mathbf{C} \boldsymbol{\sigma} \mathbf{a})} \rho_\phi(\boldsymbol{\zeta}) = -\boldsymbol{\mu}^\top \mathbf{a} + \kappa \sqrt{\mathbf{a}^\top \boldsymbol{\sigma}^\top \mathbf{C} \boldsymbol{\sigma} \mathbf{a}}.$$

Finally, since $\kappa > 0$, we have

$$\sup_{\mathbb{P} \in \mathcal{P}} \rho_\phi(-\mathbf{a}^\top \boldsymbol{\xi}) = \sup_{\mathbf{C} \in \mathcal{C}} \left\{ -\boldsymbol{\mu}^\top \mathbf{a} + \kappa \sqrt{\mathbf{a}^\top \boldsymbol{\sigma}^\top \mathbf{C} \boldsymbol{\sigma} \mathbf{a}} \right\} = -\boldsymbol{\mu}^\top \mathbf{a} + \kappa \sqrt{v(\mathbf{a})}.$$

This completes the proof. $\square$

A.3. Proof of Lemma 1

Proof. First, notice that $\mathbf{1}^\top(\gamma \mathbf{x}^*) = \gamma \mathbf{1}^\top \mathbf{x}^* = \gamma$, implying that $\gamma \mathbf{x}^* \geq \mathbf{0}$ is a feasible solution to $P(\gamma)$. We then argue the optimality by contradiction. Suppose that $\gamma \mathbf{x}^*$ is not an optimal solution to the problem $P(\gamma)$. Then, there exists $\tilde{\mathbf{x}} \geq \mathbf{0}$ such that $\mathbf{1}^\top \tilde{\mathbf{x}} = \gamma$ with objective function value strictly smaller than that of $\gamma \mathbf{x}^*$. Therefore, the observation of $-\boldsymbol{\mu}^\top \tilde{\mathbf{x}} + \kappa \sqrt{\tilde{\mathbf{x}}^\top \boldsymbol{\Sigma} \tilde{\mathbf{x}}} < -\boldsymbol{\mu}^\top(\gamma \mathbf{x}^*) + \kappa \sqrt{(\gamma \mathbf{x}^*)^\top \boldsymbol{\Sigma}(\gamma \mathbf{x}^*)}$ implies

$$-\boldsymbol{\mu}^\top \left(\frac{\tilde{\mathbf{x}}}{\gamma} \right) + \kappa \sqrt{\left(\frac{\tilde{\mathbf{x}}}{\gamma} \right)^\top \boldsymbol{\Sigma} \left(\frac{\tilde{\mathbf{x}}}{\gamma} \right)} < -\boldsymbol{\mu}^\top \mathbf{x}^* + \kappa \sqrt{(\mathbf{x}^*)^\top \boldsymbol{\Sigma} \mathbf{x}^*}.$$

Together with $\mathbf{1}^\top(\tilde{\mathbf{x}}/\gamma) = 1$, we arrive at a contradiction that $\mathbf{x}^*$ is not the optimal solution to $P(1)$.

For the second part of the lemma, consider the dual problem

$$D(\gamma) : \max_{\boldsymbol{\lambda} \geq \mathbf{0}, \pi} \min_{\mathbf{x}} \left\{ -\boldsymbol{\mu}^\top \mathbf{x} + \kappa \sqrt{\mathbf{x}^\top \boldsymbol{\Sigma} \mathbf{x}} + \pi(\gamma - \mathbf{1}^\top \mathbf{x}) - \boldsymbol{\lambda}^\top \mathbf{x} \right\}.$$

Reformulating the above problem by a constant $\gamma > 0$, we have

$$D(\gamma) : \quad \gamma \max_{\lambda \geq 0, \pi} \left\{ \pi + \min_{\mathbf{x}} \left\{ -\boldsymbol{\mu}^\top \mathbf{x} + \kappa \sqrt{\mathbf{x}^\top \boldsymbol{\Sigma} \mathbf{x}} - \pi \mathbf{1}^\top \mathbf{x} - \boldsymbol{\lambda}^\top \mathbf{x} \right\} \right\}.$$

The expression without γ is the same as problem $D(1)$, namely $D(\gamma)$ with $\gamma = 1$. Hence, the optimal solution of $D(1)$ is optimal to $D(\gamma)$ for any $\gamma > 0$. $\square$

A.4. Proof of Theorem 2

Proof. We prove by proceeding backwardly along the time horizon. Recall the loss variable at time T is $Q_T(\mathbf{x}_{T-1}, \mathbf{r}_T) = -\mathbf{r}_T^\top \mathbf{x}_{T-1}$. By Theorem 1, we have

$$\begin{aligned} Q_T(\mathbf{x}_{T-1}) &= \sup_{\mathbb{P}_T \in \mathcal{P}_T} \rho_\phi(Q_T(\mathbf{x}_{T-1}, \mathbf{r}_T)) = -\boldsymbol{\mu}^\top \mathbf{x}_{T-1} + \kappa \sqrt{\bar{v}(\mathbf{x}_{T-1})} \\ &= \sup_{\mathbf{C} \in \mathcal{C}} \left\{ -\boldsymbol{\mu}^\top \mathbf{x}_{T-1} + \kappa \sqrt{\mathbf{x}_{T-1}^\top (\boldsymbol{\sigma} \mathbf{C} \boldsymbol{\sigma}) \mathbf{x}_{T-1}} \right\}. \end{aligned}$$

Next, recall that the loss variable at time $T-1$ is

$$Q_{T-1}(\mathbf{x}_{T-2}, \mathbf{r}_{T-1}) = \inf_{\mathbf{x}_{T-1} \in \mathcal{X}_{T-1}(\mathbf{x}_{T-2}, \mathbf{r}_{T-1})} \sup_{\mathbf{C} \in \mathcal{C}} \left\{ -\boldsymbol{\mu}^\top \mathbf{x}_{T-1} + \kappa \sqrt{\mathbf{x}_{T-1}^\top (\boldsymbol{\sigma} \mathbf{C} \boldsymbol{\sigma}) \mathbf{x}_{T-1}} \right\}, \quad (1)$$

where $\mathcal{X}_{T-1}(\mathbf{x}_{T-2}, \mathbf{r}_{T-1}) = \{\mathbf{x}_{T-1} \geq \mathbf{0} \mid \mathbf{1}^\top \mathbf{x}_{T-1} = \mathbf{r}_{T-1}^\top \mathbf{x}_{T-2}\}$. Notice that the right-hand side objective function (1) in $\mathbf{x}_{T-1}$ and $\mathbf{C}$ is convex-concave on $\mathbb{R}^n \times \mathbf{S}_+^n$, where $\mathbf{S}_+^n$ is the cone of symmetric positive semidefinite matrices of size n . Also, for fixed $\mathbf{x}_{T-1}$ and $\mathbf{C}$ respectively, the function is continuous. Moreover, the set $\mathcal{C}$ is convex and compact, and the set $\mathcal{X}_{T-1}(\mathbf{x}_{T-2}, \mathbf{r}_{T-1})$ is convex. Thus, we can apply the Sion's minimax theorem to interchange the infimum and supremum operators (Sion, 1958). Therefore,

$$Q_{T-1}(\mathbf{x}_{T-2}, \mathbf{r}_{T-1}) = \sup_{\mathbf{C} \in \mathcal{C}} \inf_{\mathbf{x}_{T-1} \in \mathcal{X}_{T-1}(\mathbf{x}_{T-2}, \mathbf{r}_{T-1})} \left\{ -\boldsymbol{\mu}^\top \mathbf{x}_{T-1} + \kappa \sqrt{\mathbf{x}_{T-1}^\top (\boldsymbol{\sigma} \mathbf{C} \boldsymbol{\sigma}) \mathbf{x}_{T-1}} \right\}. \quad (2)$$

Recall that $\mathcal{X}_{T-1}(\mathbf{x}_{T-2}, \mathbf{r}_{T-1}) = \{\mathbf{x}_{T-1} \geq \mathbf{0} \mid \mathbf{x}_{T-1}^\top \mathbf{1} = \mathbf{r}_{T-1}^\top \mathbf{x}_{T-2}\}$. The minimization problem in (2) admits the form of $P(\gamma)$ with $\gamma = \mathbf{r}_{T-1}^\top \mathbf{x}_{T-2}$. In Lemma 1, we have that

$$\begin{aligned} & \inf_{\mathbf{x}_{T-1} \in \mathcal{X}_{T-1}(\mathbf{x}_{T-2}, \mathbf{r}_{T-1})} \left\{ -\boldsymbol{\mu}^\top \mathbf{x}_{T-1} + \kappa \sqrt{\mathbf{x}_{T-1}^\top (\boldsymbol{\sigma} \mathbf{C} \boldsymbol{\sigma}) \mathbf{x}_{T-1}} \right\} \\ &= \mathbf{r}_{T-1}^\top \mathbf{x}_{T-2} \cdot \inf_{\mathbf{x}_{T-1} \in \mathcal{X}'_{T-1}} \left\{ -\boldsymbol{\mu}^\top \mathbf{x}_{T-1} + \kappa \sqrt{\mathbf{x}_{T-1}^\top (\boldsymbol{\sigma} \mathbf{C} \boldsymbol{\sigma}) \mathbf{x}_{T-1}} \right\}, \end{aligned} \quad (3)$$

where $\mathcal{X}'_{T-1} = \{\mathbf{x}_{T-1} \geq \mathbf{0} \mid \mathbf{1}^\top \mathbf{x}_{T-1} = 1\}$. Notice that the inner minimization problem in (3) is a mean-standard deviation portfolio problem, which admits an explicit formula detailed in

Dentcheva and Stock (2018). That is, we have

$$\begin{aligned} Q_{T-1}(\mathbf{x}_{T-2}, \mathbf{r}_{T-1}) &= \mathbf{r}_{T-1}^\top \mathbf{x}_{T-2} \sup_{\mathbf{C} \in \mathcal{C}} \{-A_T(\mathbf{C}, \boldsymbol{\mu} + \boldsymbol{\lambda}_{T-1}^*)\} = -\mathbf{r}_{T-1}^\top \mathbf{x}_{T-2} \inf_{\mathbf{C} \in \mathcal{C}} A_T(\mathbf{C}, \boldsymbol{\mu} + \boldsymbol{\lambda}_{T-1}^*) \\ &= -\hat{A}_T^* \mathbf{r}_{T-1}^\top \mathbf{x}_{T-2}, \end{aligned}$$

with optimal solution being $\mathbf{x}_{T-1}^* = \mathbf{r}_{T-1}^\top \mathbf{x}_{T-2} [\mathbf{1}^\top (\boldsymbol{\Sigma}_T^*)^{-1} (\boldsymbol{\mu} + \boldsymbol{\lambda}_{T-1}^* - \hat{A}_T^* \mathbf{1})]^{-1} (\boldsymbol{\Sigma}_T^*)^{-1} (\boldsymbol{\mu} + \boldsymbol{\lambda}_{T-1}^* - \hat{A}_T^* \mathbf{1})$. Observe that $\mathbf{r}_{T-1}$ and $\mathbf{x}_{T-2}$ do not appear in the optimization (3) and hence, the dual variable $\boldsymbol{\lambda}_{T-1}^*$ is independent from them. This is also intuitive since the initial wealth $\mathbf{r}_{T-1}^\top \mathbf{x}_{T-2}$ should not affect the actual *investment proportion* to each asset but the *investment amount*. Note that by assumptions (a) and (b), we have $\hat{A}_T^* \geq 0$ and $\hat{A}_T^* < \infty$ respectively. Next, applying Theorem 1 again, we can write

$$\begin{aligned} Q_{T-1}(\mathbf{x}_{T-2}) &= \sup_{\mathbb{P}_{T-1} \in \mathcal{P}_{T-1}} \rho_\phi(Q_{T-1}(\mathbf{x}_{T-2}, \mathbf{r}_{T-1})) \\ &= \sup_{\mathbf{C} \in \mathcal{C}} \hat{A}_T^* \left\{ -\boldsymbol{\mu}^\top \mathbf{x}_{T-2} + \kappa \sqrt{\mathbf{x}_{T-2}^\top (\boldsymbol{\sigma} \mathbf{C} \boldsymbol{\sigma}) \mathbf{x}_{T-2}} \right\}. \end{aligned}$$

Employing minimax theorem similarly discussed previously, we obtain

$$\begin{aligned} Q_{T-2}(\mathbf{x}_{T-3}, \mathbf{r}_{T-2}) &= \inf_{\mathbf{x}_{T-2} \in \mathcal{X}_{T-2}(\mathbf{x}_{T-3}, \mathbf{r}_{T-2})} \sup_{\mathbf{C} \in \mathcal{C}} \hat{A}_T^* \left\{ -\boldsymbol{\mu}^\top \mathbf{x}_{T-2} + \kappa \sqrt{\mathbf{x}_{T-2}^\top (\boldsymbol{\sigma} \mathbf{C} \boldsymbol{\sigma}) \mathbf{x}_{T-2}} \right\} \\ &= \hat{A}_T^* \sup_{\mathbf{C} \in \mathcal{C}} \inf_{\mathbf{x}_{T-2} \in \mathcal{X}_{T-2}(\mathbf{x}_{T-3}, \mathbf{r}_{T-2})} \left\{ -\boldsymbol{\mu}^\top \mathbf{x}_{T-2} + \kappa \sqrt{\mathbf{x}_{T-2}^\top (\boldsymbol{\sigma} \mathbf{C} \boldsymbol{\sigma}) \mathbf{x}_{T-2}} \right\} \\ &= -\hat{A}_T^* \hat{A}_{T-1}^* \mathbf{r}_{T-2}^\top \mathbf{x}_{T-3}. \end{aligned}$$

Reiterating the same argument, we reach the first-stage problem as shown below:

$$\prod_{t=2}^T \hat{A}_t^* \left[\inf_{\mathbf{x}_0 \in \mathcal{X}_0} \sup_{\mathbf{C} \in \mathcal{C}} \left\{ -\boldsymbol{\mu}^\top \mathbf{x}_0 + \kappa \sqrt{\mathbf{x}_0^\top (\boldsymbol{\sigma} \mathbf{C} \boldsymbol{\sigma}) \mathbf{x}_0} \right\} \right] = -\prod_{t=1}^T \hat{A}_t^*,$$

with $\mathbf{x}_0^* = [\mathbf{1}^\top (\boldsymbol{\Sigma}_1^*)^{-1} (\boldsymbol{\mu} + \boldsymbol{\lambda}_0^* - \hat{A}_2^* \mathbf{1})]^{-1} (\boldsymbol{\Sigma}_1^*)^{-1} (\boldsymbol{\mu} + \boldsymbol{\lambda}_0^* - \hat{A}_2^* \mathbf{1})$. $\square$

A.5. Proof of Lemma 2

Proof. Define $\Phi(x) = \int_0^x \phi(p) dp$ with $x \geq 0$. Note that $\Phi(y_0) = 0$ and $\Phi(y_N) = 1$. Observe that

$$\psi(\mathbf{p}) = \sum_{i=1}^N \phi_i(\mathbf{p}) Z_i = \sum_{i=1}^N Z_i [\Phi(y_i) - \Phi(y_{i-1})] = - \sum_{i=1}^{N-1} \Phi(y_i) (Z_{i+1} - Z_i) + Z_N.$$

It suffices to show $\sum_{i=1}^{N-1} \Phi(y_i) (Z_{i+1} - Z_i)$ is convex. More specifically, as $Z_{i+1} - Z_i > 0$, we want to show $\Phi(y_i)$ is convex in p . It is evident that by linearity of y_i in p , that is, $y_i = \sum_{j=1}^i p_j$,

the convexity of $\Phi(\cdot)$ would be sufficient to verify the claim. We begin by showing one useful inequality on $\int_0^{\lambda x} \phi(p)dp$ for any $\lambda \in (0, 1)$. By a change of variable $q = p/\lambda$,

$$\int_0^{\lambda x} \phi(p)dp = \int_0^x \phi(\lambda q)\lambda dq = \lambda \int_0^x \phi(\lambda q)dq \leq \lambda \int_0^x \phi(q)dq, \quad (4)$$

where the last inequality follows from the fact that $\lambda q < q$ and non-decreasingness of ϕ . For any $x, y \geq 0$ and $\lambda \in (0, 1)$, without loss of generality, we further assume $\lambda x \leq (1 - \lambda)y$. Therefore,

$$\begin{aligned} \Phi(\lambda x + (1 - \lambda)y) &= \int_0^{\lambda x + (1 - \lambda)y} \phi(p)dp = \int_0^{\lambda x} \phi(p)dp + \int_{\lambda x}^{(1 - \lambda)y} \phi(p)dp \\ &\leq \int_0^{\lambda x} \phi(p)dp + \int_0^{(1 - \lambda)y} \phi(p)dp \\ &\leq \lambda \Phi(x) + (1 - \lambda)\Phi(y), \end{aligned}$$

where the first inequality follows from the non-negativeness of ϕ and the second inequality follows from (4). This completes the proof. $\square$

A.6. Proof of Lemma 3

Proof. Trivially, $Q_T(\mathbf{x}_{T-1}, \mathbf{r}_T) = -\mathbf{r}_T^\top \mathbf{x}_{T-1}$ is convex in $\mathbf{x}_{T-1}$. First, we claim that convexity of $Q_t(\mathbf{x}_{t-1}, \mathbf{r}_t)$ for fixed $\mathbf{r}_t$ implies $Q_t(\mathbf{x}_{t-1})$ is convex. This is a trivial use of Proposition 6.11 in Shapiro et al. (2014), with the convexity of risk measure guaranteed in Zhu and Fukushima (2009). It then suffices to show that convexity of $Q_{t+1}(\mathbf{x}_t)$ implies $Q_t(\mathbf{x}_{t-1}, \mathbf{r}_t)$ is convex in $\mathbf{x}_{t-1}$ for fixed $\mathbf{r}_t$. Following the argument in Shapiro et al. (2014), define the function

$$\psi(\mathbf{x}_t, \mathbf{x}_{t-1}, \mathbf{r}_t) = \begin{cases} Q_{t+1}(\mathbf{x}_t), & \text{if } \mathbf{1}^\top \mathbf{x}_t - \mathbf{r}_t^\top \mathbf{x}_{t-1} = 0, \\ +\infty, & \text{otherwise.} \end{cases}$$

Convexity of $Q_{t+1}(\cdot)$ implies the convexity of $\psi(\cdot, \cdot, \mathbf{r}_t)$. Hence, $Q_t(\mathbf{x}_{t-1}, \mathbf{r}_t) = \inf_{\mathbf{x}_t \in X_t} \psi(\mathbf{x}_t, \mathbf{x}_{t-1}, \mathbf{r}_t)$ is convex for fixed $\mathbf{r}_t$. Finiteness follows from the non-emptiness of $\mathcal{X}_t$. $\square$

A.7. Proof of Lemma 4

Proof. We make use of dual representation of the coherent risk measure ρ_ϕ (see, for example, Theorem 6.5 in Shapiro et al., 2014), which states that

$$\rho_\phi(Z(\mathbf{x}, \omega)) = \sup_{\zeta \in \mathcal{U}(\mathbb{P})} \sum_{i=1}^N p_i \zeta_i Z(\mathbf{x}, \omega_i)$$

for some convex compact set $\mathcal{U}(\mathbb{P})$, which is a subset of $\mathfrak{P} = \{\zeta \in \mathbb{R}^N | \zeta \geq 0, \sum_{i=1}^N p_i \zeta_i = 1\}$. Hence,

$$\begin{aligned} \psi(\mathbf{x}) &= \sup_{\mathbb{P} \in \mathcal{P}} \rho_\phi(Z(\mathbf{x}, \omega)) = \sup_{\mathbb{P} \in \mathcal{P}} \sup_{\zeta \in \mathcal{U}(\mathbb{P})} \sum_{i=1}^N p_i \zeta_i Z(\mathbf{x}, \omega_i) \geq \sup_{\zeta \in \mathcal{U}(\mathbb{P}^*)} \sum_{i=1}^N p_i^* \zeta_i Z(\mathbf{x}, \omega_i) \\ &\geq \sup_{\zeta \in \mathcal{U}(\mathbb{P}^*)} \sum_{i=1}^N p_i^* \zeta_i [Z(\bar{\mathbf{x}}, \omega_i) + z(\bar{\mathbf{x}}, \omega_i)^\top (\mathbf{x} - \bar{\mathbf{x}})] \\ &= \psi(\bar{\mathbf{x}}) + \mathbf{g}^\top (\mathbf{x} - \bar{\mathbf{x}}), \end{aligned}$$

where the first inequality comes from choosing a particular measure $\mathbb{P}^* \in \mathcal{P}$ identified as $p^* \in \mathbb{R}^N$, the worst-case distribution according to $Z(\bar{\mathbf{x}}, \omega)$, and the second one follows from subgradient inequality. The optimal weight associated to each scenario could be identified as ϕ_i^* . $\square$

A.8. Proof of Lemma 5

Proof. We prove by backward induction in time. For notational simplicity, we may suppress the superscript on iteration l . At time T , by construction,

$$\alpha_T = \tilde{\theta}_T - \tilde{\mathbf{g}}_T^\top \bar{\mathbf{x}}_{T-1} = \sum_{i=1}^N \phi_i^* \tilde{Q}_T^i - \left(- \sum_{i=1}^N \phi_i^* \mathbf{r}_T^i \right)^\top \bar{\mathbf{x}}_{T-1} = 0.$$

Notice that $\mathbf{r}_T^i \geq 0$ and $\phi_i^* \geq 0$ imply $\tilde{\mathbf{g}}_T \geq 0$. Next, assume that $\alpha_{t+1} = 0$ and $\beta_{t+1} \leq 0$. The dual problem (D) thus, could be written as

$$\begin{aligned} (D') \quad & \underset{\boldsymbol{\mu}, \pi}{\text{maximize}} && -\pi \bar{\mathbf{x}}_{t-1}^\top \mathbf{r}_t^i \\ & \text{subject to} && \sum_{l \in \mathcal{L}} \mu_l = 1, \\ & && \pi \mathbf{1} \geq - \sum_{l \in \mathcal{L}} \mu_l \beta_{t+1}^l, \\ & && \boldsymbol{\mu} \geq \text{zero}, \pi \in \mathbb{R}. \end{aligned}$$

As (P) is always feasible, strong duality holds and the optimal value of (D') is $\tilde{Q}_t^i$. Then,

$$\alpha_t = \tilde{\theta}_t - \tilde{\mathbf{g}}_t^\top \bar{\mathbf{x}}_{t-1} = \sum_{i=1}^N \phi_i^* \tilde{Q}_t^i - \left(- \sum_{i=1}^N \phi_i^* \pi_t^i \mathbf{r}_t^i \right)^\top \bar{\mathbf{x}}_{t-1} = 0.$$

Finally, observe that from the constraint, we have $\pi \mathbf{1} \geq - \sum_{l \in \mathcal{L}} \mu_l \beta_{t+1}^l \geq 0$ by the induction assumption, which implies $\pi \geq 0$. Therefore, we have $\tilde{\mathbf{g}}_t \geq 0$ and this completes the proof. $\square$

A.9. Proof of Corollary 1

Proof. Notice that (D') and the following problem

$$\begin{aligned}
 (D'') \quad & \underset{\boldsymbol{\mu}, \pi}{\text{maximize}} && -\pi \\
 & \text{subject to} && \sum_{l \in \mathcal{L}} \mu_l = 1, \\
 & && \pi \mathbf{1} \geq -\sum_{l \in \mathcal{L}} \mu_l \beta_{t+1}^l, \\
 & && \boldsymbol{\mu} \geq \mathbf{0}, \pi \in \mathbb{R},
 \end{aligned}$$

are equivalent in the sense that they have the same optimal solutions. Trivially, (D'') does not depend on i , meaning that the dual optimal solution of (D') does not depend on i . Hence, the dual optimal solutions are the same. $\square$

A.10. Proof of Lemma 6

Proof. We prove by backward induction in time. Note that $\mathcal{Q}_T(\mathbf{x}_{T-1}) = \sup_{\mathbb{P}_T \in \mathcal{P}_T} \rho_T(Q_T(\mathbf{x}_{T-1}, \mathbf{r}_T))$ with $Q_T(\mathbf{x}_{T-1}, \mathbf{r}_T)$ being evaluated exactly. Then, for any $\bar{\mathbf{x}}_{T-1} \in \text{dom}(\mathcal{Q}_T(\cdot))$, we have

$$\mathcal{Q}_T(\mathbf{x}_{T-1}) \geq \sup_{\mathbb{P}_T \in \mathcal{P}_T} \rho_T(Q_T(\bar{\mathbf{x}}_{T-1}, \mathbf{r}_T)) + \mathbf{z}_T^\top (\mathbf{x}_{T-1} - \bar{\mathbf{x}}_{T-1}),$$

for any $\mathbf{z}_T \in \partial \mathcal{Q}_T(\bar{\mathbf{x}}_{T-1})$. By Lemma 4, we know that $\beta_T^l \in \partial \mathcal{Q}_T(\bar{\mathbf{x}}_{T-1})$. Therefore, by construction,

$$\mathcal{Q}_T(\mathbf{x}_{T-1}) \geq \left[\sup_{\mathbb{P}_T \in \mathcal{P}_T} \rho_T(Q_T(\bar{\mathbf{x}}_{T-1}^l, \mathbf{r}_T)) - (\beta_T^l)^\top \bar{\mathbf{x}}_{T-1}^l \right] + (\beta_T^l)^\top \mathbf{x}_{T-1} = \alpha_T^l + (\beta_T^l)^\top \mathbf{x}_{T-1},$$

for $l = 1, 2, \dots$. Finally, we arrive at $\mathcal{Q}_T(\mathbf{x}_{T-1}) \geq \max_{l=1, \dots, k} \{\alpha_T^l + (\beta_T^l)^\top \mathbf{x}_{T-1}\} = \hat{\mathcal{Q}}_T^k(\mathbf{x}_{T-1})$.

Next, assume that $\mathcal{Q}_{t+1}(\cdot) \geq \hat{\mathcal{Q}}_{t+1}^k(\cdot)$ for all $k \in \mathbb{N}$. Notice that for $t = 1, \dots, T-1$,

$$Q_t(\mathbf{x}_{t-1}, \mathbf{r}_t) = \inf_{\mathbf{x}_t \in \mathcal{X}_t(\mathbf{x}_{t-1}, \mathbf{r}_t)} \mathcal{Q}_{t+1}(\mathbf{x}_t) \geq \inf_{\mathbf{x}_t \in \mathcal{X}_t(\mathbf{x}_{t-1}, \mathbf{r}_t)} \hat{\mathcal{Q}}_{t+1}^k(\mathbf{x}_{t-1}) = \tilde{Q}_t^k(\mathbf{x}_{t-1}, \mathbf{r}_t).$$

Therefore, for any $\bar{\mathbf{x}}_{t-1} \in \text{dom}(\sup_{\mathbb{P}_t \in \mathcal{P}_t} \rho_t(\tilde{Q}_t^k(\cdot, \mathbf{r}_t)))$,

$$\begin{aligned}
 \mathcal{Q}_t(\mathbf{x}_{t-1}) &= \sup_{\mathbb{P}_t \in \mathcal{P}_t} \rho_t(Q_t(\mathbf{x}_{t-1}, \mathbf{r}_t)) \geq \sup_{\mathbb{P}_t \in \mathcal{P}_t} \rho_t(\tilde{Q}_t^k(\mathbf{x}_{t-1}, \mathbf{r}_t)) \\
 &\geq \sup_{\mathbb{P}_t \in \mathcal{P}_t} \rho_t(\tilde{Q}_t^k(\bar{\mathbf{x}}_{t-1}^l, \mathbf{r}_t)) + (\beta_t^l)^\top (\mathbf{x}_{t-1} - \bar{\mathbf{x}}_{t-1}^l) = \alpha_t^l + (\beta_t^l)^\top \mathbf{x}_{t-1},
 \end{aligned}$$

for $l = 1, 2, \dots$, where the first inequality follows from monotonicity of the risk measure with the induction assumption and the second one comes from subgradient inequality. Taking maximum over $l = 1, \dots, k$ yields the desired result. $\square$

A.11. Proof of Lemma 7

Proof. It is easy to observe that, by construction, $\tilde{Q}_t^k(\mathbf{x}_{t-1}, \mathbf{r}_t) \geq \tilde{Q}_t^l(\mathbf{x}_{t-1}, \mathbf{r}_t)$ for $l = 1, \dots, k$. Then,

$$\tilde{\theta}_t^k = \sup_{\mathbb{P}_t \in \mathcal{P}_t} \rho_t \left(\tilde{Q}_t^k(\bar{\mathbf{x}}_{t-1}^k, \mathbf{r}_t) \right) \geq \sup_{\mathbb{P}_t \in \mathcal{P}_t} \rho_t \left(\tilde{Q}_t^l(\bar{\mathbf{x}}_{t-1}^k, \mathbf{r}_t) \right) \geq \rho_t^{\mathbb{P}^l} \left(\tilde{Q}_t^l(\bar{\mathbf{x}}_{t-1}^k, \mathbf{r}_t) \right),$$

where $\mathbb{P}^l \in \arg \max_{\mathbb{P} \in \mathcal{P}_t} \rho_t(\tilde{Q}_t^l(\bar{\mathbf{x}}_{t-1}^l, \mathbf{r}_t))$, identified as a vector $\mathbf{p}^l \in \mathbb{R}^N$. By subgradient inequality,

$$\rho_t^{\mathbb{P}^l} \left(\tilde{Q}_t^l(\bar{\mathbf{x}}_{t-1}^k, \mathbf{r}_t) \right) \geq \rho_t^{\mathbb{P}^l} \left(\tilde{Q}_t^l(\bar{\mathbf{x}}_{t-1}^l, \mathbf{r}_t) + (-\pi_t^l \mathbf{r}_t)^\top (\bar{\mathbf{x}}_{t-1}^k - \bar{\mathbf{x}}_{t-1}^l) \right),$$

where π_t^l is the dual optimal solution at iteration l with $\pi_t^{l,i}$ being the dual optimal solution associated to the realization $\mathbf{r}_t^i$. Next, making use of dual representation of risk measure, we have

$$\begin{aligned} & \rho_t^{\mathbb{P}^l} \left(\tilde{Q}_t^l(\bar{\mathbf{x}}_{t-1}^k, \mathbf{r}_t) + (-\pi_t^l \mathbf{r}_t)^\top (\bar{\mathbf{x}}_{t-1}^k - \bar{\mathbf{x}}_{t-1}^l) \right) \\ &= \max_{\zeta \in \mathcal{U}(\mathbb{P}^l)} \mathbb{E} \left[\zeta \left(\tilde{Q}_t^l(\bar{\mathbf{x}}_{t-1}^k, \mathbf{r}_t) + (-\pi_t^l \mathbf{r}_t)^\top (\bar{\mathbf{x}}_{t-1}^k - \bar{\mathbf{x}}_{t-1}^l) \right) \right] \\ &= \max_{\zeta \in \mathcal{U}(\mathbb{P}^l)} \sum_{i=1}^N p_i^l \zeta_i \left[\tilde{Q}_t^l(\bar{\mathbf{x}}_{t-1}^k, \mathbf{r}_t^i) + (-\pi_t^{l,i} \mathbf{r}_t^i)^\top (\bar{\mathbf{x}}_{t-1}^k - \bar{\mathbf{x}}_{t-1}^l) \right] \\ &\geq \sum_{i=1}^N p_i^l \zeta_i^l \left[\tilde{Q}_t^l(\bar{\mathbf{x}}_{t-1}^k, \mathbf{r}_t^i) + (-\pi_t^{l,i} \mathbf{r}_t^i)^\top (\bar{\mathbf{x}}_{t-1}^k - \bar{\mathbf{x}}_{t-1}^l) \right] \\ &= \sum_{i=1}^N p_i^l \zeta_i^l \tilde{Q}_t^l(\bar{\mathbf{x}}_{t-1}^k, \mathbf{r}_t^i) + \left(-\sum_{i=1}^N p_i^l \zeta_i^l \pi_t^{l,i} \mathbf{r}_t^i \right)^\top (\bar{\mathbf{x}}_{t-1}^k - \bar{\mathbf{x}}_{t-1}^l), \end{aligned}$$

where $\zeta^l \in \arg \max_{\zeta \in \mathcal{U}(\mathbb{P}^l)} \mathbb{E}[\zeta \tilde{Q}_t^l(\bar{\mathbf{x}}_{t-1}^l, \mathbf{r}_t)]$. Notice that

$$\rho_t^{\mathbb{P}^l} \left(\tilde{Q}_t^l(\bar{\mathbf{x}}_{t-1}^l, \mathbf{r}_t) \right) = \max_{\zeta \in \mathcal{U}(\mathbb{P}^l)} \mathbb{E} \left[\zeta \tilde{Q}_t^l(\bar{\mathbf{x}}_{t-1}^l, \mathbf{r}_t) \right] = \sum_{i=1}^N p_i^l \zeta_i^l \tilde{Q}_t^l(\bar{\mathbf{x}}_{t-1}^l, \mathbf{r}_t^i),$$

and $\beta_t^l = \sum_{i=1}^N p_i^l \zeta_i^l (-\pi_t^{l,i} \mathbf{r}_t^i)$, which are obtained in iteration l . To sum up, we arrive at the inequality

$$\tilde{\theta}_t^k = \sup_{\mathbb{P}_t \in \mathcal{P}_t} \rho_t \left(\tilde{Q}_t^k(\bar{\mathbf{x}}_{t-1}^k, \mathbf{r}_t) \right) \geq \tilde{\theta}_t^l + (\beta_t^l)^\top (\bar{\mathbf{x}}_{t-1}^k - \bar{\mathbf{x}}_{t-1}^l)$$

and therefore $\hat{Q}_t^k(\bar{\mathbf{x}}_{t-1}^k) = \max_{l=1, \dots, k} \{ \tilde{\theta}_t^l + (\beta_t^l)^\top (\bar{\mathbf{x}}_{t-1}^k - \bar{\mathbf{x}}_{t-1}^l) \} = \tilde{\theta}_t^k$. $\square$

A.12. Proof of Theorem 3

Proof. We prove by proceeding backwardly in time. Trivially, the statement in (a) holds when $t = T + 1$ as $Q_{T+1} = 0$ and $\hat{Q}_T^k = 0$ by construction. Now, assume that the statement holds for

time $t + 1$, i.e. $\lim_{k \rightarrow \infty} [\mathcal{Q}_{t+1}(\mathbf{x}_t^k) - \widehat{\mathcal{Q}}_{t+1}^k(\mathbf{x}_t^k)] = 0$. For a scenario $\tilde{\mathbf{r}}_{t-1}$ at time $t - 1$, define the set

$$\mathcal{S}_{\tilde{\mathbf{r}}_{t-1}} = \{k \in \mathbb{N} \mid \bar{\mathbf{r}}_{t-1} = \tilde{\mathbf{r}}_{t-1}\}$$

as the iterations such that the node $\tilde{\mathbf{r}}_{t-1}$ is selected. By Assumption 2, we have $\mathcal{S}_{\tilde{\mathbf{r}}_{t-1}}$ contains infinite elements. We continue by showing the convergence under two cases: $k \in \mathcal{S}_{\tilde{\mathbf{r}}_{t-1}}$ and $k \notin \mathcal{S}_{\tilde{\mathbf{r}}_{t-1}}$.

First, consider the case when $k \in \mathcal{S}_{\tilde{\mathbf{r}}_{t-1}}$. Then, the decision associated to that node is $\bar{\mathbf{x}}_{t-1}^k$, the solution obtained in the forward pass. Note that due to Lemmas 6 and 7,

$$\begin{aligned} 0 &\leq \mathcal{Q}_t(\bar{\mathbf{x}}_{t-1}^k) - \widehat{\mathcal{Q}}_t^k(\bar{\mathbf{x}}_{t-1}^k) = \sup_{\mathbb{P}_t \in \mathcal{P}_t} \rho_t \left(\mathcal{Q}_t(\bar{\mathbf{x}}_{t-1}^k, \mathbf{r}_t) \right) - \sup_{\mathbb{P}_t \in \mathcal{P}_t} \rho_t \left(\widetilde{\mathcal{Q}}_t^k(\bar{\mathbf{x}}_{t-1}^k, \mathbf{r}_t) \right) \\ &= \sup_{\mathbb{P}_t \in \mathcal{P}_t} \sup_{\zeta \in \mathcal{U}(\mathbb{P})} \mathbb{E}_{\mathbb{P}} \left[\zeta \mathcal{Q}_t(\bar{\mathbf{x}}_{t-1}^k, \mathbf{r}_t) \right] - \sup_{\mathbb{P}_t \in \mathcal{P}_t} \sup_{\zeta \in \mathcal{U}(\mathbb{P})} \mathbb{E}_{\mathbb{P}} \left[\zeta \widetilde{\mathcal{Q}}_t^k(\bar{\mathbf{x}}_{t-1}^k, \mathbf{r}_t) \right]. \end{aligned}$$

For any given $\mathbb{P} \in \mathcal{P}_t$, let $\zeta^*(\mathbb{P}) \in \arg \max_{\zeta \in \mathcal{U}(\mathbb{P})} \mathbb{E}_{\mathbb{P}}[\zeta \mathcal{Q}_t(\bar{\mathbf{x}}_{t-1}^k, \mathbf{r}_t)]$. Obviously, by construction, we have $\sup_{\zeta \in \mathcal{U}(\mathbb{P})} \mathbb{E}_{\mathbb{P}}[\zeta \widetilde{\mathcal{Q}}_t^k(\bar{\mathbf{x}}_{t-1}^k, \mathbf{r}_t)] \geq \mathbb{E}_{\mathbb{P}}[\zeta^*(\mathbb{P}) \widetilde{\mathcal{Q}}_t^k(\bar{\mathbf{x}}_{t-1}^k, \mathbf{r}_t)]$, which implies

$$\begin{aligned} 0 &\leq \mathcal{Q}_t(\bar{\mathbf{x}}_{t-1}^k) - \widehat{\mathcal{Q}}_t^k(\bar{\mathbf{x}}_{t-1}^k) \\ &\leq \sup_{\mathbb{P}_t \in \mathcal{P}_t} \sum_{i=1}^N p_i \zeta_i^*(\mathbb{P}) \mathcal{Q}_t(\bar{\mathbf{x}}_{t-1}^k, \mathbf{r}_t^i) - \sup_{\mathbb{P}_t \in \mathcal{P}_t} \sum_{i=1}^N p_i \zeta_i^*(\mathbb{P}) \widetilde{\mathcal{Q}}_t^k(\bar{\mathbf{x}}_{t-1}^k, \mathbf{r}_t^i) \\ &\leq \sup_{\mathbb{P}_t \in \mathcal{P}_t} \sum_{i=1}^N p_i \zeta_i^*(\mathbb{P}) \left[\mathcal{Q}_t(\bar{\mathbf{x}}_{t-1}^k, \mathbf{r}_t^i) - \widetilde{\mathcal{Q}}_t^k(\bar{\mathbf{x}}_{t-1}^k, \mathbf{r}_t^i) \right]. \end{aligned}$$

Notice that in backward pass, we have $\widetilde{\mathcal{Q}}_t^k(\bar{\mathbf{x}}_{t-1}^k, \mathbf{r}_t^i) = f_t(\mathbf{x}_t^{k,i}, \mathbf{r}_t^i) + \widehat{\mathcal{Q}}_{t+1}^k(\mathbf{x}_t^{k,i})$ as defined, where $\mathbf{x}_t^{k,i} \in \arg \min_{\mathbf{x}_t \in \mathcal{X}_t(\bar{\mathbf{x}}_{t-1}, \mathbf{r}_t^i)} \{f_t(\mathbf{x}_t, \mathbf{r}_t^i) + \widehat{\mathcal{Q}}_{t+1}^k(\mathbf{x}_t)\}$. Hence, we arrive at

$$0 \leq \mathcal{Q}_t(\bar{\mathbf{x}}_{t-1}^k) - \widehat{\mathcal{Q}}_t^k(\bar{\mathbf{x}}_{t-1}^k) \leq \sup_{\mathbb{P}_t \in \mathcal{P}_t} \sum_{i=1}^N p_i \zeta_i^*(\mathbb{P}) \left[\mathcal{Q}_t(\bar{\mathbf{x}}_{t-1}^k, \mathbf{r}_t^i) - \widehat{\mathcal{Q}}_{t+1}^k(\mathbf{x}_t^{k,i}) - f_t(\mathbf{x}_t^{k,i}, \mathbf{r}_t^i) \right].$$

Finally, from the dynamic programming equation and the feasibility of $\mathbf{x}_t^{k,i}$, $\mathcal{Q}_t(\bar{\mathbf{x}}_{t-1}^k, \mathbf{r}_t^i) \leq f_t(\mathbf{x}_t^{k,i}, \mathbf{r}_t^i) + \mathcal{Q}_{t+1}(\mathbf{x}_t^{k,i})$ for all i . This gives

$$0 \leq \mathcal{Q}_t(\bar{\mathbf{x}}_{t-1}^k) - \widehat{\mathcal{Q}}_t^k(\bar{\mathbf{x}}_{t-1}^k) \leq \sup_{\mathbb{P}_t \in \mathcal{P}_t} \sum_{i=1}^N p_i \zeta_i^*(\mathbb{P}) \left[\mathcal{Q}_{t+1}(\mathbf{x}_t^{k,i}) - \widehat{\mathcal{Q}}_{t+1}^k(\mathbf{x}_t^{k,i}) \right].$$

By induction assumption, as $k \rightarrow \infty$, right hand side goes to zero. Therefore,

$$\lim_{k \rightarrow \infty, k \in \mathcal{S}_{\tilde{\mathbf{r}}_{t-1}}} \left[\mathcal{Q}_t(\bar{\mathbf{x}}_{t-1}^k) - \widehat{\mathcal{Q}}_t^k(\bar{\mathbf{x}}_{t-1}^k) \right] = 0.$$

Next, we are left with the case when $k \notin \mathcal{S}_{\tilde{\mathbf{r}}_{t-1}}$. This part adopts the idea from Girardeau et al. (2014) and Guigues (2016). We argue by contradiction. Suppose that

$$\lim_{k \rightarrow \infty, k \notin \mathcal{S}_{\tilde{\mathbf{r}}_{t-1}}} \left[\mathcal{Q}_t(\mathbf{x}_{t-1}^k) - \hat{\mathcal{Q}}_t^k(\mathbf{x}_{t-1}^k) \right] = 0$$

does not hold. Hence, there exists $\epsilon > 0$ such that the set $\mathcal{K}_{t-1,\epsilon} = \{k \in \mathbb{N} | \mathcal{Q}_t(\mathbf{x}_{t-1}^k) - \hat{\mathcal{Q}}_{t-1}^k(\mathbf{x}_{t-1}^k) \geq \epsilon\}$ is infinite. Define the random variables $w_{t-1}^k = \mathbb{1}(k \in \mathcal{K}_{t-1,\epsilon})$ and $y_{t-1}^k = \mathbb{1}(k \in \mathcal{S}_{\tilde{\mathbf{r}}_{t-1}})$, where $\mathbb{1}$ denotes the indicator function. Let z^j be the j -th element in $\{y_{t-1}^k | k \in \mathcal{K}_{t-1,\epsilon}\}$. Using Lemma A3 in Girardeau et al. (2014), we have $\{z^j\}$ being independent and identically distributed as y_{t-1}^1 . By strong law of large numbers (SLLN), as $K \rightarrow \infty$, $\bar{\mathbb{P}}$ -almost surely,

$$\frac{1}{K} \sum_{j=1}^K z^j \rightarrow \mathbb{E}_{\bar{\mathbb{P}}}(z^1) = \mathbb{E}_{\bar{\mathbb{P}}}(y_{t-1}^1) = \bar{\mathbb{P}}(\tilde{\mathbf{r}}_{t-1} = \tilde{\mathbf{r}}_{t-1}) > 0.$$

However, using the result from the first case, we have $\mathcal{K}_{t-1,\epsilon} \cap \mathcal{S}_{\tilde{\mathbf{r}}_{t-1}}$ being finite and hence $K^{-1} \sum_{j=1}^K z^j \rightarrow 0$, which gives the desired contradiction.

For part (b), notice that $\underline{z}^k = \hat{\mathcal{Q}}_1^k(\mathbf{x}_0^k)$ and $z^* \leq \mathcal{Q}_1(\mathbf{x}_0^k)$. Hence, $0 \leq z^* - \underline{z}^k \leq \mathcal{Q}_1(\mathbf{x}_0^k) - \hat{\mathcal{Q}}_1^k(\mathbf{x}_0^k)$ with the right hand side of the inequality goes to zero by part (a). This yields $\lim_{k \rightarrow \infty} \underline{z}^k = z^*$. Finally, consider the sequence $\{\mathbf{x}_0^k\}$. By the compactness of the feasibility set $\mathcal{X}_0$, there exists a convergent subsequence, denoted as $\{\mathbf{x}_0^k\}_{k \in \mathcal{K}}$ for some index set $\mathcal{K}$, with accumulation point $\mathbf{x}_0^* \in \mathcal{X}_0$. Observe that by (a) and continuity of $\mathcal{Q}_1$, we have

$$\lim_{k \rightarrow \infty, k \in \mathcal{K}} \hat{\mathcal{Q}}_1^k(\mathbf{x}_0^k) = \lim_{k \rightarrow \infty, k \in \mathcal{K}} \mathcal{Q}_1(\mathbf{x}_0^k) = \mathcal{Q}_1(\mathbf{x}_0^*).$$

Taking limit along $k \in \mathcal{K}$ on both sides of $\underline{z}^k = \hat{\mathcal{Q}}_1^k(\mathbf{x}_0^k)$, we obtain $\mathcal{Q}_1(\mathbf{x}_0^*) = z^*$. Hence, $\mathbf{x}_0^*$ is an optimal solution to the problem. $\square$

B. l_2 approximation of spectral risk measure

In this section, we detail the l_2 approximation of spectral risk measure introduced in Su et al. (2019). The idea is to search for a step function that decently approximates the original spectral function $\phi(p)$. Assume that the number of steps of the approximated spectral function $\hat{\phi}$ is n . For the l_2 approximation, we want to perform the following minimization problem.

$$\begin{aligned} \text{Minimize} \quad & \sum_{i=1}^n \int_{\alpha_{i-1}}^{\alpha_i} [\phi(p) - h_i]^2 dp \\ (P_{l_2}) \quad \text{subject to} \quad & \sum_{i=1}^n h_i(\alpha_i - \alpha_{i-1}) = 1 \\ & \alpha_{i-1} \leq \alpha_i \quad , \quad i = 1, 2, \dots, n, \end{aligned}$$

where $\alpha_0 = 0$ and $\alpha_n = 1$. To perform minimization over α and h together is difficult whereas for (P_{l_2}) , it can be reformulated to a simpler problem (see Theorem 4 in [Su et al., 2019](#)). We first present a generalized version of the theorem and apply the result to our case.

Theorem B.1. *Let $f : [a, b] \rightarrow \mathbb{R}$ be a non-decreasing function. For the minimization problem $\min_h \int_a^b [f(x) - h]^2 dx$, the optimal solution is $h^* = \frac{1}{b-a} \int_a^b f(x) dx \in [f(a), f(b)]$.*

Proof. Notice that the optimal solution must lie in $[f(a), f(b)]$. Write the objective function as

$$g(h) := \int_a^b [f(x) - h]^2 dx = \int_a^b [f(x)]^2 dx - 2h \int_a^b f(x) dx + h^2(b-a),$$

which is a quadratic function in h . We remark here that $f \in L_2[a, b]$ to ensure the finiteness of the first term. It is elementary to see that the minima is attained at $h^* = \frac{1}{b-a} \int_a^b f(x) dx$. $\square$

Consider the l_2 cost function in (P_{l_2}) . Making use of Theorem B.1 on individual integrals, we have $h_i^* = \frac{1}{\alpha_i - \alpha_{i-1}} \int_{\alpha_{i-1}}^{\alpha_i} \phi(p) dp$ being the minimizers. Interestingly, the first constraint would be satisfied automatically for such a choice of h . Indeed,

$$\sum_{i=1}^n h_i^* (\alpha_i - \alpha_{i-1}) = \sum_{i=1}^n \int_{\alpha_{i-1}}^{\alpha_i} \phi(p) dp = 1,$$

implying that the optimal h for whichever choice of α must take the form of h^* . That is, the problem is equivalent to the following one:

$$(P'_{l_2}) \quad \begin{array}{ll} \text{Minimize} & \sum_{i=1}^n \int_{\alpha_{i-1}}^{\alpha_i} [\phi(p) - h_i^*(\alpha_{i-1}, \alpha_i)]^2 dp \\ \text{subject to} & \alpha_{i-1} \leq \alpha_i \quad , \quad i = 1, 2, \dots, n. \end{array}$$

By removing the constant term $\int_0^1 [\phi(p)]^2 dp$ in the cost function (assuming that $\phi \in L_2[0, 1]$), we have

$$(P_{l_2}) \quad \begin{array}{ll} \text{Minimize} & - \sum_{i=1}^n \frac{1}{\alpha_i - \alpha_{i-1}} \left[\int_{\alpha_{i-1}}^{\alpha_i} \phi(p) dp \right]^2 \\ \text{subject to} & \alpha_{i-1} \leq \alpha_i \quad , \quad i = 1, 2, \dots, n. \end{array}$$

As we assume that $\Phi(x) = \int_0^x \phi(p) dp$ is defined on $[0, 1]$, the term $\int_{\alpha_{i-1}}^{\alpha_i} \phi(p) dp$ can be replaced by $\Phi(\alpha_i) - \Phi(\alpha_{i-1})$. Notice that this problem is non-convex in general even in the constrained domain (see Figure 1 for an example).

A general approach for numerical optimization problem is gradient methods. For simplicity, we assume that $\Phi(x)$ is differentiable with derivative $\phi(x)$ on $(0, 1)$. It would suffice to

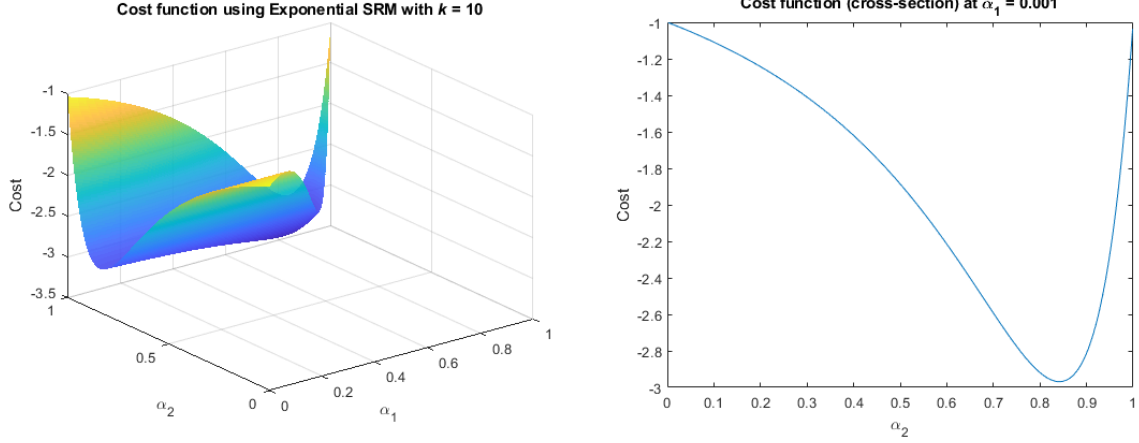


Figure 1: Non-convexity of l_2 cost with $\alpha_1 \leq \alpha_2$ ($n = 3$) on exponential type with $k = 10$

incorporate the popular spectral functions in the literature. Elementary calculus gives the formula for the gradient for the cost function, which is denoted as $J(\alpha)$.

$$\frac{\partial J}{\partial \alpha_i} = (s_i^2 - s_{i+1}^2) - 2\phi(\alpha_i) [s_i - s_{i+1}], \text{ where } s_i = \frac{\Phi(\alpha_i) - \Phi(\alpha_{i-1})}{\alpha_i - \alpha_{i-1}} \quad (5)$$

for $i = 1, 2, \dots, n-1$. To perform the minimization problem, a ‘gradient projection algorithm’ is provided in [Su et al. \(2019\)](#). Let $\theta^{(k)}$ be the step size at iteration k , which is diminishing, ϵ be the tolerance and K be the maximum number of iterations. (Indeed, a constant step size or backtracking step size is possible.) The algorithm they proposed is detailed in [Algorithm 1](#).

Algorithm 1: Approximating spectral function by l_2 loss in [Su et al. \(2019\)](#)

Initialization: Initialize $\alpha_1^{(0)}, \dots, \alpha_{n-1}^{(0)}$ with $\alpha_{i-1} \leq \alpha_i$, $i = 1, 2, \dots, n$, $k = 1$

Iterate until no significant change in variables or iteration number exceeds pre-specified value (ignore the while condition for the 1st iteration)

while $\|\alpha^{(k)} - \alpha^{(k-1)}\| \geq \epsilon$ **and** $k \leq K$ **do**

(Gradient Method)

 Compute the gradients $\frac{\partial J}{\partial \alpha_i}$ for $i = 1, \dots, n-1$ using [\(5\)](#).

 Set $\alpha^{(k)} = \alpha^{(k-1)} - \theta^{(k)} \frac{\partial J}{\partial \alpha}$ and sort the values of $\alpha^{(k)}$ in ascending order.

 Set $k = k + 1$.

end

It should be noticed that the algorithm does not guarantee the generated $\alpha^{(k)}$ must lie

between 0 and 1. A proper choice of step size is, therefore, a crucial practical problem. Also, this algorithm differs from classical gradient projection method (a type of proximal gradient method) that the projection is conducted by the minimization problem

$$\begin{aligned} & \text{Minimize} \quad -\frac{1}{2} \|\mathbf{u} - \boldsymbol{\alpha}\|_2^2 \\ & \text{subject to} \quad \alpha_{i-1} \leq \alpha_i \quad , \quad i = 1, 2, \dots, n, \end{aligned} \tag{6}$$

where $\mathbf{u}$ is the point generated. This would be a simple quadratic programming problem and could be solved efficiently by interior-point methods. The algorithm is thus modified as in Algorithm 2.

Algorithm 2: Approximating spectral function by l_2 loss by gradient proejction method

Initialization: Initialize $\alpha_1^{(0)}, \dots, \alpha_{n-1}^{(0)}$ with $\alpha_{i-1} \leq \alpha_i$, $i = 1, 2, \dots, n$, $k = 1$

Iterate until no significant change in variables or iteration number exceeds pre-specified value (ignore the while condition for the 1st iteration)

while $\|\boldsymbol{\alpha}^{(k)} - \boldsymbol{\alpha}^{(k-1)}\| \geq \epsilon$ **and** $k \leq K$ **do**

(Gradient Projection Method)

Compute the gradients $\frac{\partial J}{\partial \alpha_i}$ for $i = 1, \dots, n-1$ using (5).

Set $\hat{\boldsymbol{\alpha}} = \boldsymbol{\alpha}^{(k-1)} - \theta^{(k)} \frac{\partial J}{\partial \boldsymbol{\alpha}}$ and solve (6) with $\mathbf{u} = \hat{\boldsymbol{\alpha}}$ for $\boldsymbol{\alpha}^{(k)}$.

Set $k = k + 1$.

end

Remark B..1. Notice that if the values of two consecutive α_i are the same, the computation of the gradient may not be possible as both numerator and denominator are zero. For the minimization problem (6), the use of interior-point method may help avoid some practical problems during iteration as optimal solution on the boundary are slightly inclined to the interior part. This is demonstrated and discussed in the following example.

Example B..1. Consider the approximation of exponential spectral function with parameter $k = 10$ using 3 steps, which means we have two variables α_1 and α_2 . The step size is chosen as $\theta^{(k)} = 0.1/\sqrt{k}$, which is diminishing and non-summable. The above two algorithms are performed with $\epsilon = 10^{-7}$. The objective function value throughout the iterations and the contour plot are shown in Figures 2 and 3, respectively.

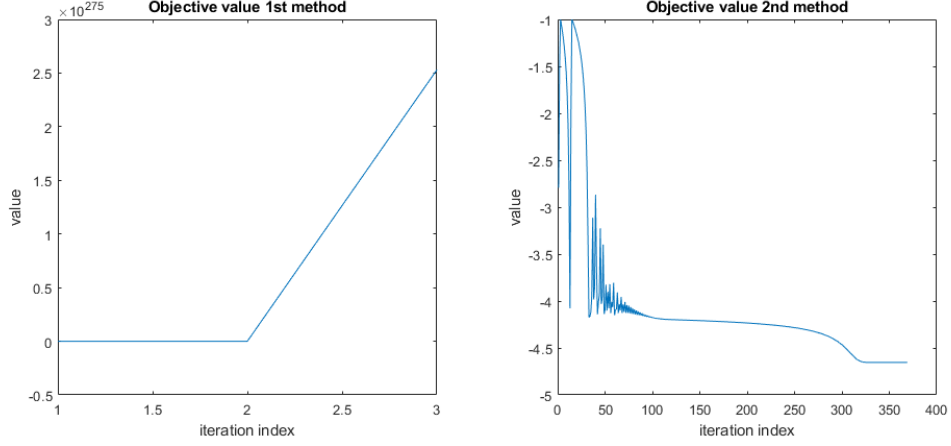


Figure 2: Objective function value over iterations via Algorithm 1 (left) and 2 (right)

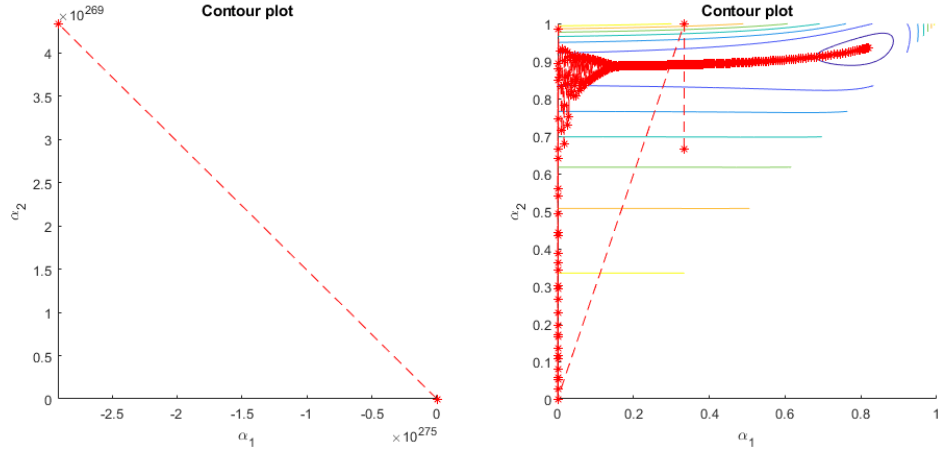


Figure 3: Contour plot and the iteration points via Algorithm 1 (left) and 2 (right)

The points generated using Algorithm 1 after the first few iterations lie far outside the feasibility region and the objective function value explodes with the iteration stopped after 3 iterations. While using the gradient projection method, the point converges to the minima using 368 iterations. One should notice that the step size should not be too large as it would slow down the convergence rate, which is a typical issue in numerical optimization methods. In the case that the steps are sufficiently small that the iterative points would lie outside the feasibility region, they perform similarly. This implies that Algorithm 2 permits flexibility in the choice of step size.

Finally, we remark on the use of interior-point methods. Notice that from Figure 3 (right),

the 3rd iteration point is very close to $(0,0)$. Indeed, the exact solution to (6) is $(0,0)$. If the exact solution is used, the gradient in the next step would be undefined. With the use of interior-point methods, the approximated solution avoids this practical problem.

C. Additional computational results on simulation study

C.1. Additional computational results for Section 6.1

In Section 6.1, we conduct the experiment using 45 stocks and examine the convergence property of our numerical algorithm. In this section, based on the estimated mean, standard deviation and correlation of the 45 assets, we conduct a simulation study to examine the effect on the number of scenarios.

First, based on the summary statistics of the assets, we generate 50 scenarios using multivariate normal distribution. Then, we construct the ambiguity set based on the estimated mean and standard deviation of the 50 scenarios, as well as the confidence intervals of the correlations. Then, we simulate extra $N - 50$ scenarios (so that they add up to N scenarios) and solve the problem using the proposed numerical method. As in Section 6.1, we use the approximated exponential type spectral risk measure with parameter $h \in \{2, 10\}$. We consider various number of periods $T \in \{1, 2, 5\}$. Figures 4–6 show the convergence under different number of scenarios $N \in \{200, 400, 800, 1600\}$ for each pair of (T, h) . The dashed line indicates the optimal value from the semi-analytical solution. We observe that when we increase the number of scenarios, the stabilized lower bound keeps increasing. Moreover, the improvement of the objective value is diminishing. That is, there is a relatively larger improvement when we double the number of scenarios from $N = 200$ than that from $N = 400$. Nevertheless, we observe that all the values are below the optimal value from the semi-analytical solution. Based on the empirical $1/N$ convergence of the simulation instances we observed, we also compute the approximated limit of the converging sequence in Table 1. The extrapolation values are very close to the SDP optimal solution. Note that we use multivariate normal distribution for the simulation study, where the support is the $\mathbb{R}^n$ space. Therefore, we could expect that if when N is very large, we could obtain an optimal value close the one computed from the semi-analytical solution.

Next, we also study the effect on convergence using different initial guesses. We repeat the same experiment by directly generating $N \in \{200, 500\}$ scenarios using the estimated

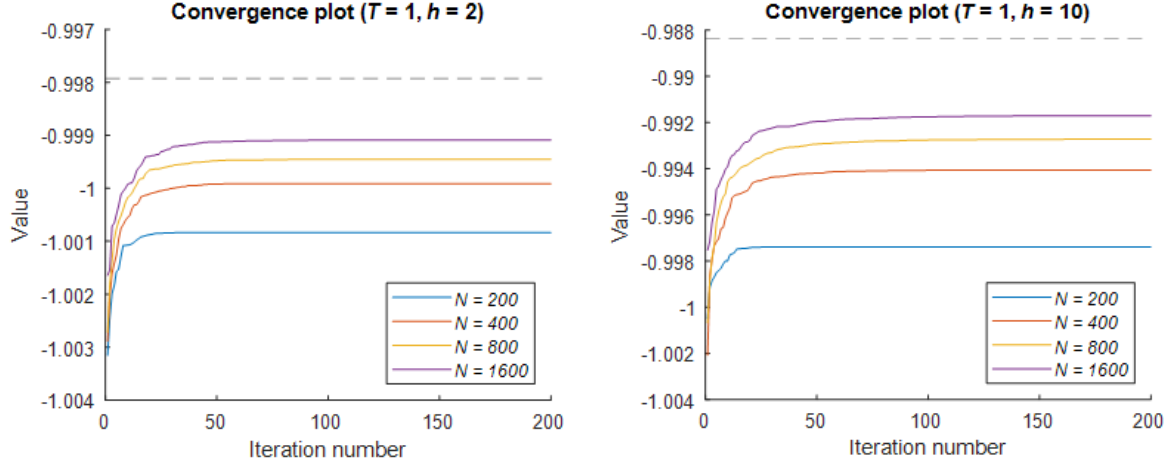


Figure 4: Performance under different number of scenarios when $T = 1$

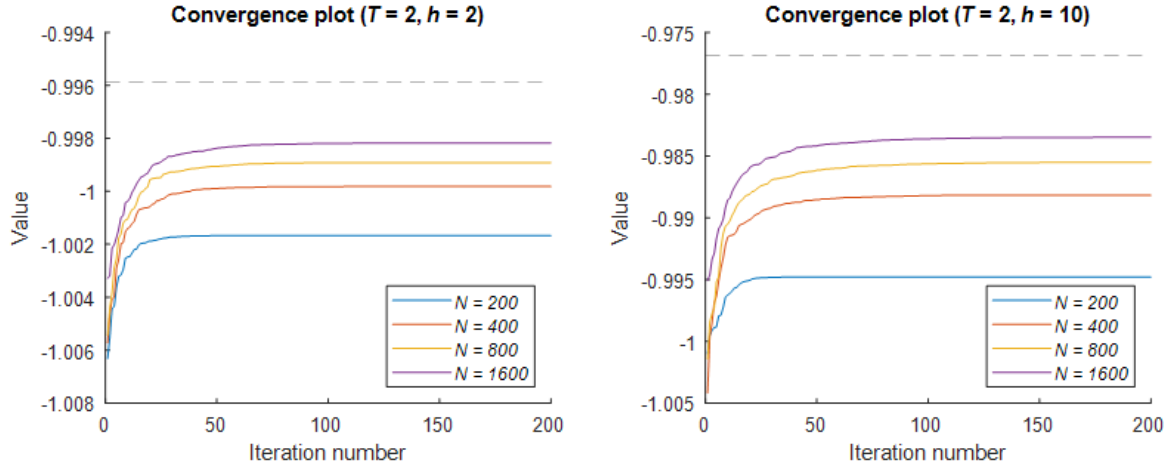


Figure 5: Performance under different number of scenarios when $T = 2$

N	200	400	800	1600	Richardson	SDP
$(T = 1, h = 2)$	-1.0008	-0.9999	-0.9995	-0.9991	-0.9986	-0.9979
$(T = 1, h = 10)$	-0.9974	-0.9941	-0.9927	-0.9917	-0.9903	-0.9884
$(T = 2, h = 2)$	-1.0017	-0.9998	-0.9989	-0.9982	-0.9971	-0.9959
$(T = 2, h = 10)$	-0.9948	-0.9882	-0.9855	-0.9835	-0.9807	-0.9769
$(T = 5, h = 2)$	-1.0042	-0.9995	-0.9973	-0.9955	-0.9929	-0.9897
$(T = 5, h = 10)$	-0.9870	-0.9707	-0.9643	-0.9594	-0.9526	-0.9432

Table 1: Optimal value and its Richardson extrapolation

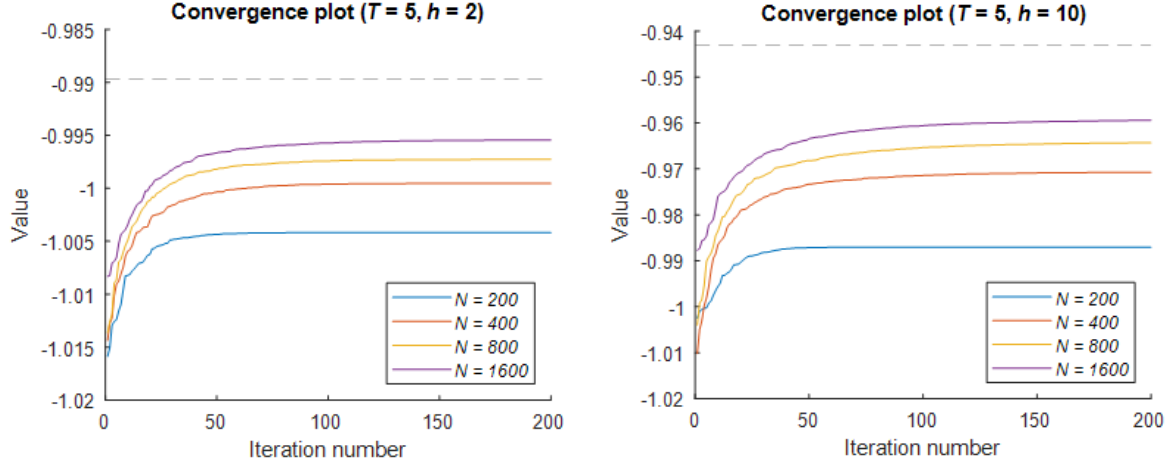


Figure 6: Performance under different number of scenarios when $T = 5$

parameters via multivariate normal distribution. We consider three different initial guesses: (a) the equally-weighted (EW) portfolio, i.e., $(1/45, \dots, 1/45)^\top \in \mathbb{R}^{45}$, (b) the optimal solution from semi-analytical (SA) solution and (c) a uniformly generated vector on $(0, 1)$ after normalization (such that sum of the weights is one). Since the observations are similar, we only present the results when $T = 5$ and $h = 10$ in Figure 7. While there are some minor differences in the lower bound values at the beginning, there are no significant differences in terms of convergence. That is, the lower bounds from different initial guesses stabilize after a certain number of iterations. This demonstrates that our proposed method is relatively insensitive to the initial guess.

C.2. Simulation setting and results for Section 6.3

In Section 6.3, we consider two settings for scenario generations. We generate scenarios with multivariate normal distribution, with the first one being estimated based on historical data of five stocks in HSI composition and the second one obtained by changing the mean. Table 2 gives the precise values for simulation, denoted as setting I and setting II.

The plot, similar to Figure 5 in the main text, under setting II is shown in the left plot of Figure 8. We also test for the case when the second batch of correlation structure is generated in an arbitrary way. For the arbitrary correlation structure, the following correlation matrix is

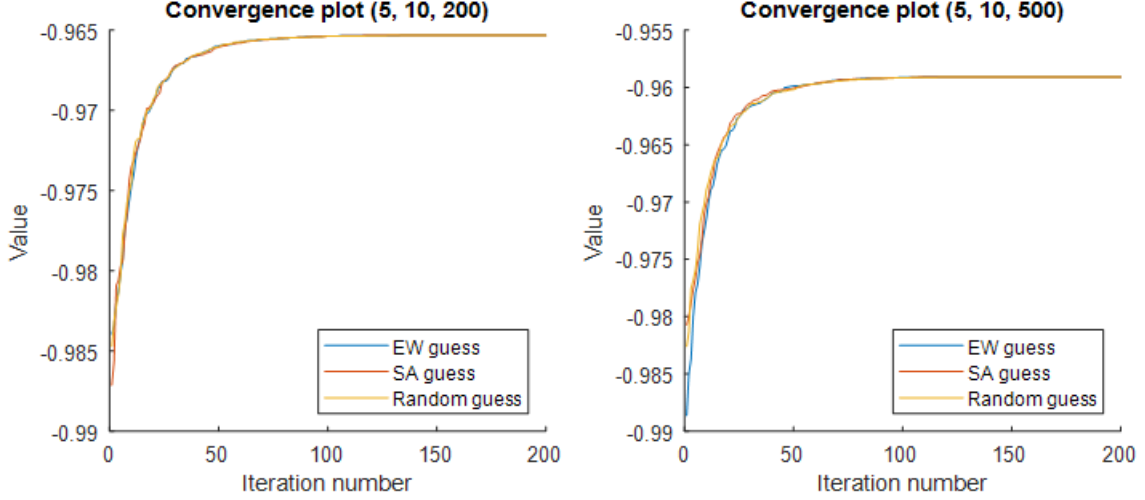


Figure 7: Convergence with different initial guesses under $T = 5$ and $h = 10$ (Left: $N = 200$, Right: $N = 500$)

employed with most of the coefficients being selected out of the correlation ambiguity bounds.

$$\begin{pmatrix} 1 & 0.30 & 0.50 & 0.41 & 0.31 \\ 0.30 & 1 & 0.42 & 0.43 & 0.51 \\ 0.50 & 0.42 & 1 & 0.32 & 0.46 \\ 0.41 & 0.43 & 0.32 & 1 & 0.22 \\ 0.31 & 0.51 & 0.46 & 0.22 & 1 \end{pmatrix}$$

The plot is shown in the right plot of Figure 8. All the figures shown behave analogously as in Figure 5. The simulation study reveals the importance of correlation ambiguity in the portfolio model.

D. Supplementary computational results on real data analysis

D.1. Other choices of parameter r in Section 6.3

We conduct the same numerical test in Section 6.3 with some other choices of r , namely $r = 0.01$ and $r = 0.03$. The parameter r corresponds to the radius of the l_2 uncertainty set in Philpott et al. (2018). We present the summary statistics for different models in Table 3. We can observe that the model with $r = 0.03$ shares a similar pattern with the one with $r = 0.05$, which is discussed in section 6.3. Among the three choices of r , the one with $r = 0.01$ performs the

	Mean	Standard Dev.	Correlation
Setting I	$\begin{pmatrix} 1.000405 \\ 1.000224 \\ 1.000351 \\ 1.000155 \\ 1.000216 \end{pmatrix}$	$\begin{pmatrix} 0.013910 \\ 0.008220 \\ 0.009936 \\ 0.011776 \\ 0.011481 \end{pmatrix}$	$\begin{pmatrix} 1 & 0.425 & 0.401 & 0.468 & 0.343 \\ 0.425 & 1 & 0.554 & 0.335 & 0.468 \\ 0.401 & 0.554 & 1 & 0.404 & 0.461 \\ 0.468 & 0.335 & 0.404 & 1 & 0.291 \\ 0.343 & 0.468 & 0.461 & 0.291 & 1 \end{pmatrix}$
Setting II	$\begin{pmatrix} 0.999595 \\ 0.999776 \\ 0.999649 \\ 0.999845 \\ 0.999784 \end{pmatrix}$	$\begin{pmatrix} 0.013910 \\ 0.008220 \\ 0.009936 \\ 0.011776 \\ 0.011481 \end{pmatrix}$	$\begin{pmatrix} 1 & 0.425 & 0.401 & 0.468 & 0.343 \\ 0.425 & 1 & 0.554 & 0.335 & 0.468 \\ 0.401 & 0.554 & 1 & 0.404 & 0.461 \\ 0.468 & 0.335 & 0.404 & 1 & 0.291 \\ 0.343 & 0.468 & 0.461 & 0.291 & 1 \end{pmatrix}$

Table 2: Settings for scenario generation in section 7.2

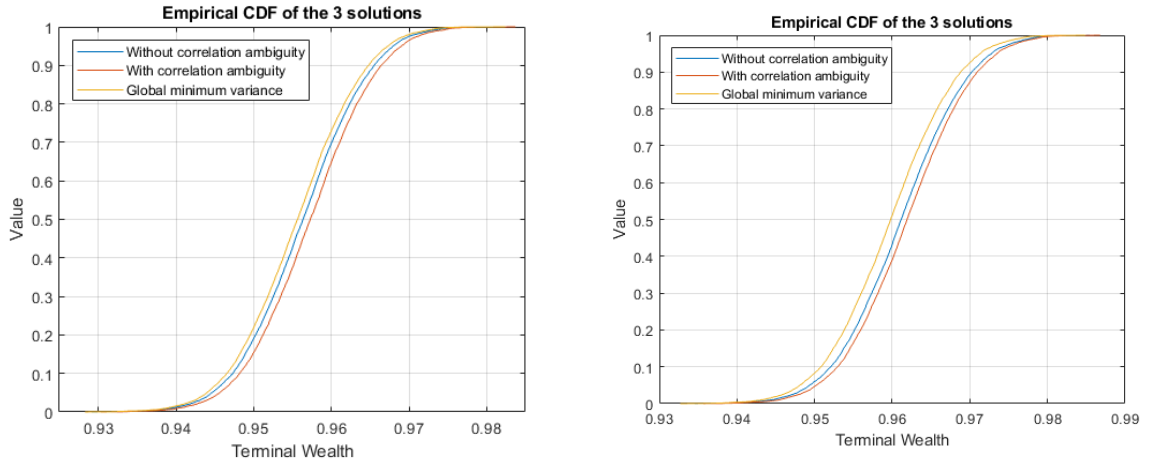


Figure 8: Simulation study under setting II (left) and the one with an arbitrary second batch correlation (right)

best, mainly due to its fast recovery after the crisis (see Figure 9). This leads to a relatively larger variance overall though the mean-to-sd ration outperforms all other examined models.

Model	Mean ($\times 10^{-4}$)	Variance ($\times 10^{-4}$)	Mean-to-SD
Proposed model with $h = 2$	1.7467	1.1270	0.01645
Proposed model with $h = 10$	2.0333	1.1041	0.01935
Proposed model with $h = 15$	2.0545	1.0852	0.01972
Global min. variance model	1.6542	0.9480	0.01699
DRO Exp model with $r = 0.01$	2.1145	1.1202	0.01998
DRO Exp model with $r = 0.03$	1.9630	1.1018	0.01870
DRO Exp model with $r = 0.05$	1.9471	1.0774	0.01876

Table 3: Out-of-sample daily return statistics

However, we note that the proposed model with $h = 15$ eventually attains a terminal wealth similar to the DRO expectation model with $r = 0.01$.

D.2. Performance in a longer period

In this section, we provide results for an analysis similar to section 6.3 performed using a longer period, which starts from 01/07/2003 to 16/06/2019. The economic downturn at the end of 2008 is also incorporated, followed by an upward trend. We first compare the performance of the three choices of h and then compare them with the other two models mentioned in section 6.3. The procedure is the same as the one before.

Figure 10 shows the result for the whole period. During the shorter period from 2007 to 2013, we observe that the model with $h = 15$ performs the best. However, from this long-period result, we observe that the model with $h = 10$ outperforms the other two eventually. We may explain by first inspecting the wealth processes before the financial crisis. Due to the stable upward trend during the period of 2003 to 2007, there is a salient distinction between the three models. The one with $h = 2$ outshines the other two, which is ascribed to the smallest degree of risk averseness, while the one with $h = 15$ displays the reverse situation. This could be observed clearly in the lower panel. The differences reduce during the financial crisis and its succeeding recovery period. During the uptrend period, the model with $h = 10$ becomes the dominating one since the model with $h = 15$ is more risk-averse and the model with $h = 2$ invests more in BOVESPA, which has a large mean but poor overall performance during the period around

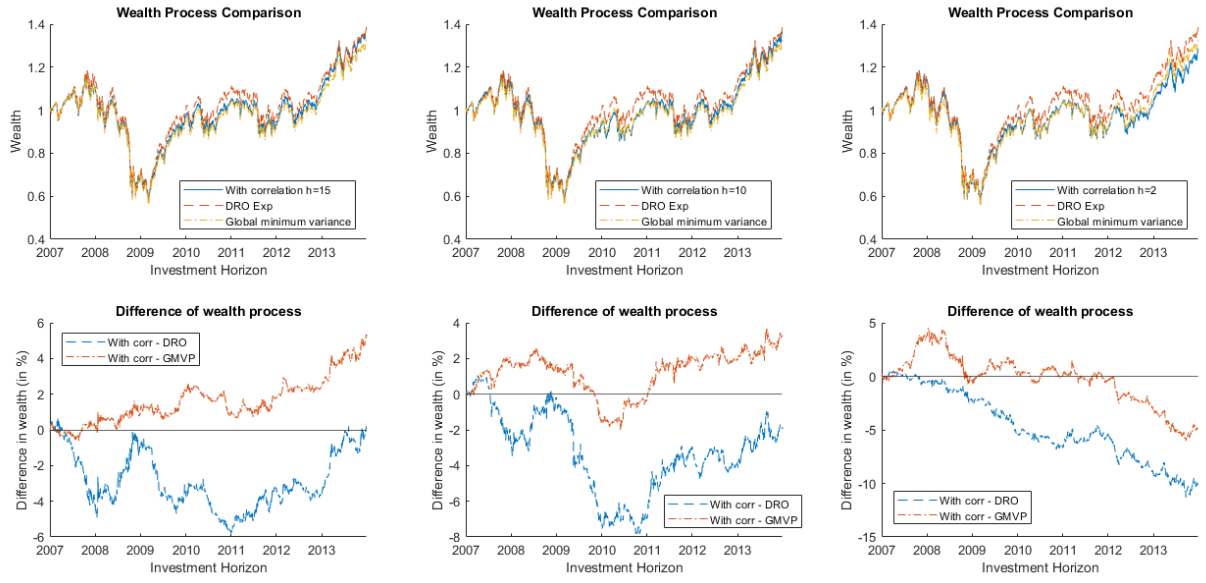


Figure 9: Wealth process and difference comparing with different models: GMVP and DRO Exp with $r = 0.01$

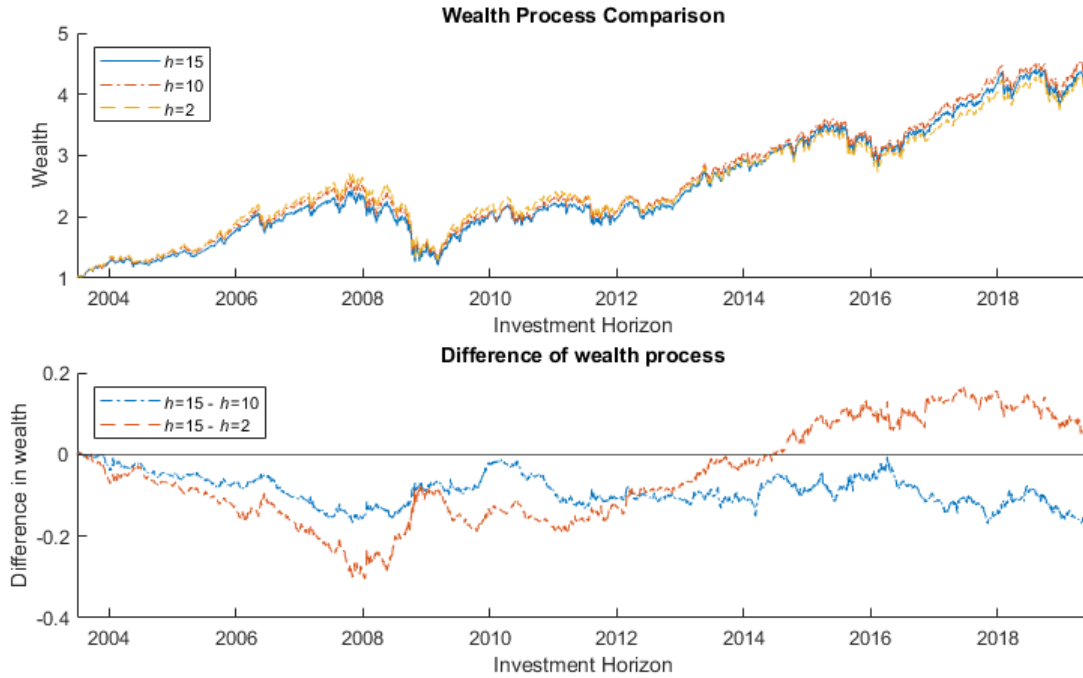


Figure 10: Wealth process and difference for different values of h

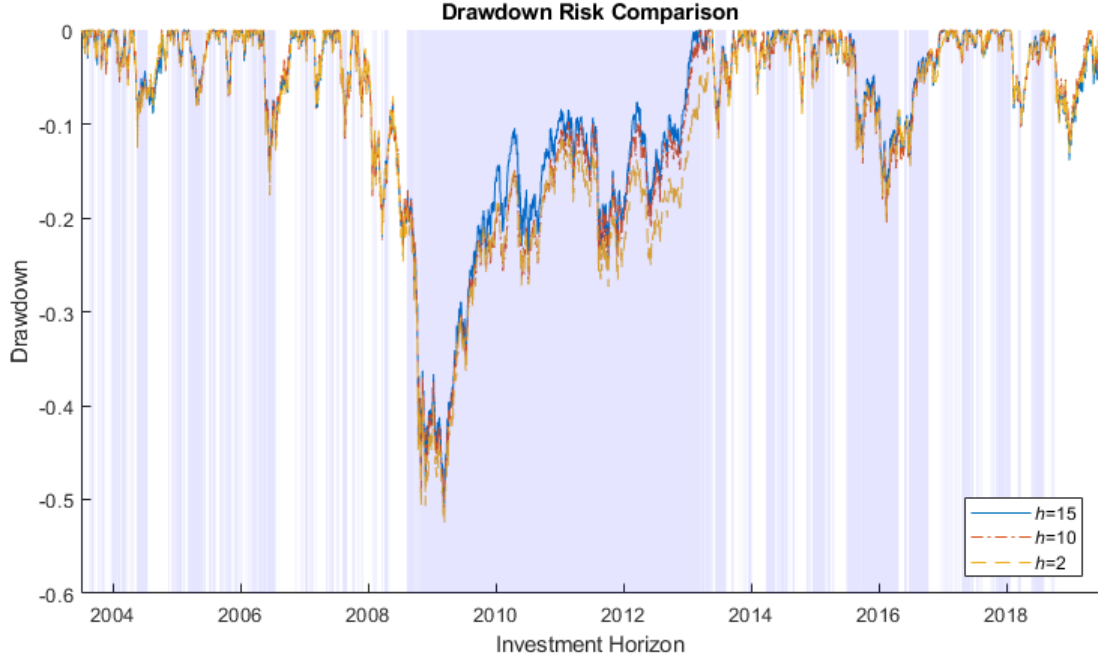


Figure 11: Drawdown of the model with three choices of h

2014 to 2015. The small downturn in 2016 leads to the catch up of the difference between the models $h = 10$ and $h = 15$ but by then the stable upward trend leads to the preeminence of the model with $h = 10$.

Figure 11 demonstrates the drawdown of the three wealth processes and the regions colored represent the time when the model $h = 15$ has the least drawdown. A similar pattern as in Figure 7 could be observed. During the financial crisis around 2008 and the subsequent recovery period, the drawdown of the model with $h = 15$ is the smallest. Indeed, this also emerges at time of 2016. This demonstrates the protective power induced from the risk averseness of our model, although it may not have the best terminal wealth, especially in the periods including a continuous positive trend.

Finally, we compare the three choices of h in our proposed model with the other two models specified in section 6.3. We first discuss the case $h = 10$ which the overall terminal wealth is the largest. We observe that it outperforms the single-period global minimum variance model to a remarkable extent, which is more than 10%. The difference between the second one, which is also a robust model, stays positive most of the time. For $h = 15$, we observe that the performance is relatively poor in the beginning of the testing period although the differences are ameliorated during the downturn periods, which leads to the underachieving terminal wealth.

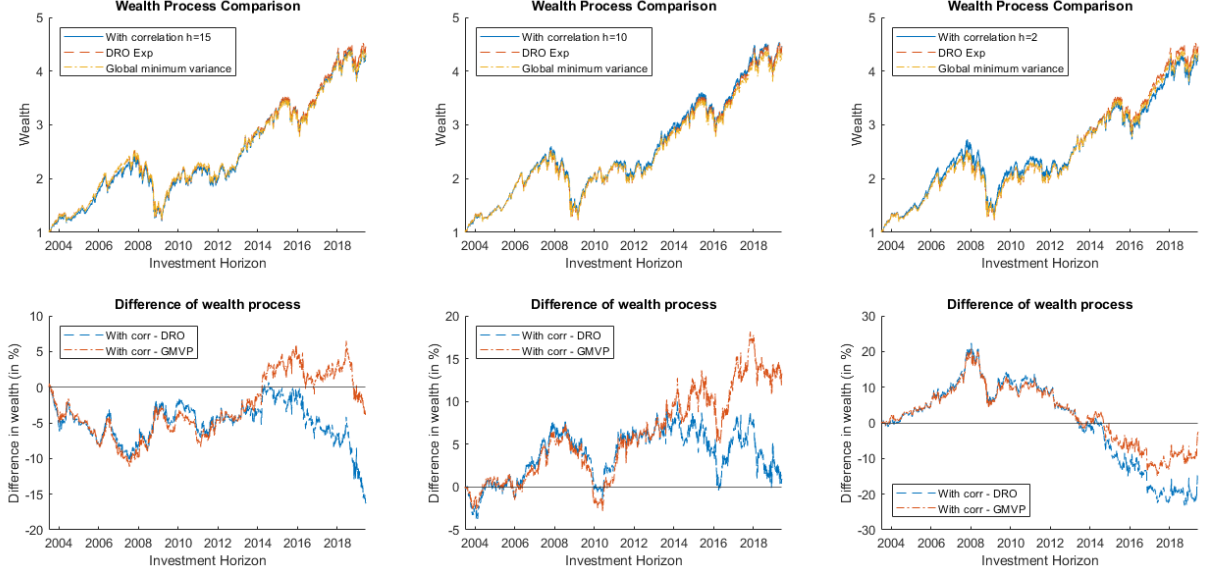


Figure 12: Wealth process and difference comparing with different models

The behaviour is quite disparate when $h = 2$. In the beginning, it performs better than the other two models due to the aggressiveness. This argument also holds when compared to the one using l_2 norm for ambiguity set as our model requires the distribution having the same mean extracted from historical data. However, the advantage dissipates in the later periods, leading to a significant large negative differences ultimately.

We conclude that using our model with a medium level of risk averseness would be beneficial due to a compromise between financial market downturn and uptrend. However, as described in Section 6.3, if a downturn is anticipated in the near future, higher level of risk aversion should yield a better terminal wealth. The choice of the parameters, which alters the degree of risk aversion, is beyond the scope of this paper, but it is suggested that a change in parameter in different time periods may further improve the return.

E. Notation summary

Indices	
n	Number of assets
N	Number of scenarios
T	Number of stages
Parameters	
α	Confidence level for CVaR
$\mathbf{\Gamma}$	Variance of a random vector $\boldsymbol{\xi}$ in $\mathbb{R}^n$
$\boldsymbol{\lambda}_t$	Dual variable associated to the constraint $\mathbf{x}_t \geq 0$
μ_i	Mean of the asset i with $\boldsymbol{\mu} = (\mu_1, \dots, \mu_n)^\top$
σ_i^2	Variance of the asset i with $\boldsymbol{\sigma} = \text{diag}(\sigma_1, \dots, \sigma_n)$
Σ	Variance matrix of the n assets
$\mathbf{C}_\pi$	Correlation matrix under distribution π
$\overline{\mathbf{C}}$	Upper bound for the correlation matrix
$\underline{\mathbf{C}}$	Lower bound for the correlation matrix
h	Parameter for the exponential spectral risk measure
$\mathbf{p}$	Probability vector $(p_1, \dots, p_N)$
$\mathbf{r}_t$	One plus the arithmetic return at time t ($\mathbf{r}_t \in \mathbb{R}^n$)
W_t	Wealth at time t with $W_0 = 1$
$\mathbf{x}_t$	Investment decision at time t ($\mathbf{x}_t \in \mathbb{R}^n$)
Sets	
$\mathcal{C}$	Set of correlation matrices which are lower and upper bounded by $\underline{\mathbf{C}}$ and $\overline{\mathbf{C}}$
$\mathcal{D}$	Set of all probability distributions
$\mathcal{F}_t$	Filtration up to time t
Ω_t	Sample space at time t , which is $\{\omega_t^1, \dots, \omega_t^N\}$ under finite scenario setting
$\mathcal{P}_t$	Uncertainty set at time t (we also use $\mathcal{P}$ for simplicity)
$\mathcal{P}_+$	Set of distributions in $\mathcal{P}$ with non-negative support
$\mathcal{Z}$	Set of random variables, which we treat it as the space of L^2 random variables
Functions	
ρ_t	Risk functional employed at time t conditional on $\mathcal{F}_{t-1}$
ρ_ϕ	Spectral risk measure with spectral function ϕ
f_t	Cost function at time t
F_Z	Distribution function of a random variable Z
F_Z^{-1}	Left-continuous quantile function of a random variable Z
Q_t	Dynamic equation at time t
$\mathcal{X}_t$	Feasible set at time t , which is a multifunction
Notation related to the numerical method	
α_t^k	Intercept coefficient of the cut for Q_t approximation at iteration k
β_t^k	Slope coefficient of the cut for Q_t approximation at iteration k
ϵ	Lower bound absolute tolerance
π_t^i	Dual optimal solution associated to the equality constraint in (13) at the i -th scenario
$\boldsymbol{\mu}_t^i$	Dual optimal solution associated to the inequality constraint in (13) at the i -th scenario
K	Maximum number of iterations
$\bar{\mathcal{P}}$	Sampling distribution in the forward pass
Q_t	Recourse /Cost-to-go function at time t
$\hat{Q}_t$	Lower linear approximation of Q_t
$\hat{Q}_t^k$	Approximation of Q_t after k iterations
$\tilde{Q}_t^k$	Dynamic programming equation by replacing Q_t by $\tilde{Q}_t^k$
$\bar{\mathbf{x}}_t$	Trial point of the time t problem in forward passes
$\mathcal{X}_t$	Feasible set at time t , which is a multifunction
$\underline{z}^k$	Approximated objective value at iteration k
z^*	Optimal solution of the discretized problem

References

- Chen, L., He, S., and Zhang, S. (2011), “Tight bounds for some risk measures, with applications to robust portfolio selection,” *Oper. Res.*, 59, 847–865.
- Dentcheva, D. and Stock, G. J. (2018), “On the price of risk in a mean-risk optimization model,” *Quantitative Finance*, 18, 1699–1713.
- Girardeau, P., Leclere, V., and Philpott, A. B. (2014), “On the convergence of decomposition methods for multistage stochastic convex programs,” *Math. Oper. Res.*, 40, 130–145.
- Guigues, V. (2016), “Convergence analysis of sampling-based decomposition methods for risk-averse multistage stochastic convex programs,” *SIAM J. Optim.*, 26, 2468–2494.
- Kuhn, D., Mohajerin Esfahani, P., Nguyen, V. A., and Shafieezadeh-Abadeh, S. (2019), “Wasserstein distributionally robust optimization: Theory and applications in machine learning,” in *Operations Research & Management Science in the Age of Analytics*, INFORMS, pp. 130–166.
- Li, J. Y.-M. (2018), “Closed-form solutions for worst-case law invariant risk measures with application to robust portfolio optimization,” *Oper. Res.*, 66, 1533–1541.
- Philpott, A. B., de Matos, V. L., and Kapelevich, L. (2018), “Distributionally robust SDDP,” *Comput. Manag. Sci.*, 15, 431–454.
- Shapiro, A., Dentcheva, D., and Ruszczyński, A. (2014), *Lectures on Stochastic Programming: Modeling and Theory*, SIAM, 2nd ed.
- Sion, M. (1958), “On general minimax theorems,” *Pacific Journal of mathematics*, 8, 171–176.
- Su, L., Sun, L., Karwan, M., and Kwon, C. (2019), “Spectral risk measure minimization in hazardous materials transportation,” *IIE Trans.*, 1–15.
- Zhu, S. and Fukushima, M. (2009), “Worst-case conditional value-at-risk with application to robust portfolio management,” *Oper. Res.*, 57, 1155–1168.