
Robust Decision-Focused Learning via Worst-Case Regret Minimization

Shoki Yamao¹ Ken Kobayashi¹ Ryo Matsui² Shota Nagai¹ Naoki Nishimura² Kazuhide Nakata¹

¹Institute of Science Tokyo, Tokyo, Japan

²Recruit Co., Ltd., Chiyoda, Tokyo, Japan,

Abstract

In optimization-based decision-making, when the objective coefficient vector is unknown, a common approach is to predict it from covariates using a predictive model and solve the downstream optimization problem. However, improving predictive accuracy does not necessarily lead to better decisions. This has motivated Decision-Focused Learning (DFL), which trains predictive models to minimize decision loss measured by regret. Despite recent progress, existing DFL methods do not sufficiently address two sources of uncertainty: (i) observation errors in the measured coefficients, and (ii) distribution shift in the coefficient vector at deployment due to environmental changes. These uncertainties can degrade solution quality. In this paper, we propose two robust regret losses to address these uncertainties. For uncertainty (i), we introduce an uncertainty set around the observed coefficient vector that captures measurement errors and define the loss as the worst-case regret over this set. For uncertainty (ii), we construct a Wasserstein ambiguity set around the conditional empirical distribution and define the loss using the worst-case distribution within the set. We enable efficient training via Danskin-based subgradients. Experiments demonstrate that our method reduces regret and yields more stable solutions than existing robust DFL approaches.

1 INTRODUCTION

Many real-world decision-making tasks can be formulated as mathematical optimization problems, including shortest-path and resource-allocation problems [Hillier and Lieberman, 2015, Bertsimas and Tsitsiklis, 1997]. In such problems, parameters such as travel times and demands appear as coefficients in the objective function. We denote this co-

efficient vector by $c \in \mathbb{R}^n$. However, these coefficients are often unknown at the time of decision-making, so the optimization problem cannot be solved directly. A widely studied approach is the two-stage predict-then-optimize pipeline: given covariates $x \in \mathcal{X}$ (e.g., weather, seasonality, and historical demand), a predictive model f_θ estimates the unknown coefficients as $\hat{c} = f_\theta(x)$, and we then solve the downstream optimization problem

$$z(\hat{c}) := \min_{w \in \mathcal{W}} \hat{c}^\top w, \quad (1)$$

where $\mathcal{W}$ is a nonempty compact feasible set, w denotes the decision variables, and $\theta \in \Theta$ denotes the parameters of the predictive model [Bertsimas and Kallus, 2020, Huber et al., 2019, Rios et al., 2015, El Balghiti et al., 2023, Elmachtoub and Grigas, 2022, Sadana et al., 2025]. As a concrete example, let us consider energy scheduling, where the coefficients c are the hourly electricity prices. A model produces predicted prices $\hat{c}$ from covariates x , and these prices are used to determine an hourly consumption schedule w subject to operational constraints. Since the realized prices are revealed only after the schedule is committed, the decision must be made before c is known.

Prior work typically trained the predictive model f_θ by minimizing standard prediction losses such as mean squared error. Under this paradigm, the goal is to make the predicted coefficient vector $\hat{c}$ match the ground-truth coefficients c as closely as possible, with the expectation that this would improve the quality of the downstream optimal solution [Bertsimas and Kallus, 2020, Huber et al., 2019, Rios et al., 2015]. However, it has been recognized that higher predictive accuracy does not necessarily lead to better downstream solutions [Elmachtoub and Grigas, 2022, Sadana et al., 2025, Donti et al., 2017, Wilder et al., 2019].

Motivated by this observation, recent work has moved away from focusing solely on matching the predicted coefficients to the ground-truth coefficients. Instead, these approaches directly target the quality of the solution obtained when solving the downstream optimization problem with those

predictions [Elmachtoub and Grigas, 2022, Sadana et al., 2025, Donti et al., 2017, Wilder et al., 2019, Kotary et al., 2021, Mandi et al., 2024, Bengio et al., 2021]. This paradigm is known as Decision-Focused Learning (DFL). DFL aims to couple prediction and optimization by training model parameters to directly minimize a decision loss (e.g., regret), thereby improving downstream solution quality under the predicted coefficient vector.

Despite this progress, existing DFL methods often do not adequately account for the following practical uncertainties in real-world deployments:

- **Uncertainty (i):** Observation errors caused by sensor aging in training data.
- **Uncertainty (ii):** Distribution shift between training and deployment, where the conditional distribution of the cost coefficients given the covariates changes due to environmental variations.

Because conventional approaches train predictive models without explicitly accounting for (i) and (ii), the resulting decisions can be brittle, and the downstream solution quality can degrade substantially. This motivates a new DFL framework that incorporates these uncertainties to enable robust decision-making.

Prior work has incorporated ideas from robust optimization into DFL to handle Uncertainty (i) [Schutte et al., 2024, Wang et al., 2026]. Schutte et al. [2024] study robustness to uncertainty arising from limited data and noisy observations and propose three robust loss formulations. Wang et al. [2026] use a generative model to learn the distribution of coefficient vectors, aiming to learn a model that is robust to the predicted distribution.

However, important limitations remain. First, existing robust DFL methods [Schutte et al., 2024, Wang et al., 2026] do not explicitly control the worst-case decision loss induced by observation errors, leaving them vulnerable to noisy measurements. Second, a robust DFL for Uncertainty (ii) has not been established where the underlying data-generating process changes between training and deployment. Most approaches minimize empirical risk under the training distribution, offering no guarantees when the testing distribution deviates. We discuss these limitations in more detail in Section 2.

To address these issues, we propose a unified framework that evaluates the worst-case regret under uncertainty. Based on this framework, we introduce two robust loss functions targeting each source of uncertainty:

- **For Uncertainty (i) (Observation Errors):** Robust Regret Loss $\mathcal{L}_{\text{rob}}$, which minimizes the worst-case regret over a parameter-based uncertainty set.
- **For Uncertainty (ii) (Distribution Shift):** Distributionally Robust Regret Loss $\mathcal{L}_{\text{dro}}$, which minimizes the worst-case expected regret over a Wasserstein am-

biguity set.

Concretely, we leverage robust optimization techniques [Ben-Tal et al., 2009, Bertsimas and Sim, 2004] to model observation errors (Uncertainty (i)) and distributionally robust optimization (DRO) techniques [Delage and Ye, 2010, Mohajerin Esfahani and Kuhn, 2018] to model distribution shifts (Uncertainty (ii)), yielding uncertainty sets for the coefficient vector. We then formulate a learning objective that minimizes the worst-case decision loss over these sets. Our goal is to obtain models that reliably produce high-quality solutions even in the presence of observation errors or distribution shifts in the coefficient vector.

However, directly minimizing these worst-case measures presents computational challenges because the optimization involves infinite-dimensional constraints or nested maximization problems, making direct gradient computation difficult. To enable efficient training, we reformulate these optimization problems into tractable forms solvable by standard techniques and derive subgradients for the proposed losses.

The main contributions of this paper are as follows:

1. To address observation errors, we develop a robust DFL loss that models such errors via a parameter-space uncertainty set and minimizes the worst-case decision loss.
2. To handle distribution shifts, we propose a distributionally robust DFL loss that minimizes the expected decision loss under the worst-case conditional distribution within a Wasserstein ambiguity set.
3. We reduce the infinite-dimensional distribution optimization problem to a computationally tractable regularized loss function.
4. We derive subgradients for the proposed worst-case loss functions to establish an efficient gradient-based learning algorithm.
5. Numerical experiments demonstrate that our method improves downstream solution quality under uncertainty in terms of average regret and risk metrics compared with existing DFL baselines.

2 PRELIMINARIES

2.1 DECISION-FOCUSED LEARNING

In approaches that decouple learning a prediction model $f_{\theta} : \mathcal{X} \rightarrow \mathbb{R}^n$ from solving the downstream optimization problem, the training loss typically targets prediction accuracy. However, more accurate predictions of the coefficient vector $\hat{c} = f_{\theta}(x)$ do not necessarily yield better decisions, i.e., better solutions $w^*(\hat{c}) \in \arg \min_{w \in \mathcal{W}} \hat{c}^{\top} w$. Decision-focused learning (DFL), in contrast, trains f_{θ} end-to-end to improve downstream decision quality by incorporating the optimization problem into the learning objective.

The downstream decision quality is commonly quantified by the regret [Sadana et al., 2025]:

$$\text{Reg}(\hat{c}, c) := c^\top w^*(\hat{c}) - z(c). \quad (2)$$

This quantity measures how much the decision optimized for the prediction, $w^*(\hat{c})$, degrades the true objective value relative to the true optimum under c .

In standard DFL, given training data $\mathcal{D} = \{(x_i, c_i)\}_{i=1}^N$, we learn the model parameters θ by minimizing the empirical average regret:

$$\min_{\theta} \frac{1}{N} \sum_{i=1}^N \text{Reg}(f_{\theta}(x_i), c_i). \quad (3)$$

Because $\text{Reg}(\hat{c}, c)$ is generally nonconvex and discontinuous in the prediction $\hat{c}$, practitioners typically optimize a continuous convex surrogate, such as the SPO+ loss [El-machtoub and Grigas, 2022].

2.2 RELATED WORK

Due to space limitations, we provide a brief overview of related work here. A more detailed discussion is available in Appendix H.

2.2.1 Existing research on Uncertainty (i)

Several approaches have been proposed to address Uncertainty (i) in DFL. Schutte et al. [2024] proposed three robust loss functions to handle Uncertainty (i). The first method, the k -NN loss, replaces the coefficient vector c in (2) with a local average $\bar{c}$ computed from neighbors in the covariate space. However, this approach can be sensitive to large observation errors and does not guarantee worst-case regret robustness. The second method is the RO loss, which uses a robust solution $w^{\text{RO}}(c)$ defined with respect to an uncertainty set $\mathcal{C}$:

$$\mathcal{L}_{\text{RO}}(\hat{c}, c) = c^\top w^*(\hat{c}) - c^\top w^{\text{RO}}(c). \quad (4)$$

Here, $w^{\text{RO}}(c) \in \arg \min_{w \in \mathcal{W}} \max_{\delta \in \mathcal{C}} (c + \delta)^\top w$. Crucially, the second term $c^\top w^{\text{RO}}(c)$ in (4) does not depend on the prediction $\hat{c}$. Consequently, minimizing the RO loss is equivalent to minimizing the standard regret up to a constant shift. The third method also minimizes the standard regret up to a constant shift, failing to explicitly control the worst-case regret induced by Uncertainty (i).

Gen-DFL [Wang et al., 2026] is another approach addressing Uncertainty (i). It trains a generative model q to capture conditional distributions and minimizes the CVaR of regret. However, if the generative model q is trained using observed coefficient vectors c_{obs} that contain errors, it may learn the observation-noise structure itself. If the learned distribution deviates from the true underlying distribution, optimizing against q does not guarantee robustness under the true distribution.

2.2.2 Existing research on Uncertainty (ii)

While DRO is well-studied, to the best of our knowledge, there is no existing work that applies it to the DFL setting to address Uncertainty (ii). As noted in Uncertainty (ii), standard DFL objectives do not account for shifts in the conditional distribution. Without explicit distributional robustness, the optimizer $w^*(\hat{c})$ can be fragile to out-of-distribution data.

3 PROPOSED METHOD

In this section, we present a framework for robust decision-focused learning based on worst-case regret minimization. We study worst-case regret minimization under two uncertainty models: a set-based model for observation errors and a Wasserstein distributional model for deployment shifts. Although these uncertainties arise from different sources, we address them through a common principle: minimizing the worst-case decision loss over a defined uncertainty set. Sections 3.1 and 3.2 detail the specific formulations for observation errors and distribution shifts, respectively.

3.1 DECISION-FOCUSED LEARNING ROBUST TO OBSERVATION ERRORS

In this subsection, we propose a robust DFL loss that accounts for observation errors in the coefficient vector. We first define a robust regret loss function based on worst-case regret minimization over an uncertainty set. Next, we derive its subgradient for model training. However, as the computational cost of the subgradient is high, we also introduce an approximate robust regret loss to improve learning efficiency.

3.1.1 Robust Regret Loss Function

To address observation errors, we regard the observed coefficient vector c_{obs} as a noisy observation of the true, unobserved coefficient vector c , and model the range of plausible perturbations by an uncertainty set $\mathcal{C}(c_{\text{obs}})$ centered at c_{obs} . Our goal is to learn parameters θ such that the induced decision $w^*(\hat{c})$, with $\hat{c} = f_{\theta}(x)$, remains reliable against any coefficient vector c within this uncertainty set. Accordingly, we evaluate robustness by the worst-case regret over $\mathcal{C}(c_{\text{obs}})$ and define the following robust regret loss:

$$\mathcal{L}_{\text{rob}}(\hat{c}, c_{\text{obs}}) := \max_{c \in \mathcal{C}(c_{\text{obs}})} \{c^\top w^*(\hat{c}) - z(c)\}, \quad (5)$$

where the uncertainty set $\mathcal{C}(c_{\text{obs}})$ is nonempty and compact.

This loss corresponds to the regret incurred under the worst-case coefficient vector $c \in \mathcal{C}(c_{\text{obs}})$. Training to minimize $\mathcal{L}_{\text{rob}}$, we obtain a model whose decisions remain robust to observation errors, i.e., it reduces worst-case regret and limits the degradation of solution quality. Typical choices for $\mathcal{C}(c_{\text{obs}})$ include ellipsoidal and budgeted uncertainty

sets [Ben-Tal et al., 2009, Bertsimas and Sim, 2004]; their definitions and detailed formulations are provided in Appendix B.1.

3.1.2 Gradient Calculation of Robust Regret Loss

To train a predictive model (e.g., a neural network) with the robust regret loss (5), we need to compute gradients of the loss with respect to the model output $\hat{c}$ and update the parameters θ . However, the robust regret loss $\mathcal{L}_{\text{rob}}$ involves a maximization over the uncertainty set, and its objective further contains the optimal value function $z(c)$. This nested structure makes gradient computation nontrivial. In this section, we derive a subgradient of $\mathcal{L}_{\text{rob}}$ via Danskin’s theorem. We then obtain a practical training algorithm by combining the resulting subgradient with the surrogate gradient.

Using Danskin’s theorem [Bertsekas, 1999], we derive a subgradient of the robust regret loss $\mathcal{L}_{\text{rob}}$ with respect to the decision $w^*(\hat{c})$ obtained from the predicted coefficient vector $\hat{c}$. The subdifferential is characterized by the following theorem. (See Appendix D.1 for the proof.)

Theorem 3.1. *Let $\mathcal{W}$ and $\mathcal{C}$ be nonempty compact sets. Then the subdifferential of $\mathcal{L}_{\text{rob}}(\hat{c}, c_{\text{obs}})$ with respect to $w^*(\hat{c})$ is given by the convex hull of the worst-case coefficient vectors, i.e.,*

$$\partial_{w^*(\hat{c})} \mathcal{L}_{\text{rob}}(\hat{c}, c_{\text{obs}}) = \text{conv} \left\{ \arg \max_{c \in \mathcal{C}(c_{\text{obs}})} (c^\top w^*(\hat{c}) - z(c)) \right\}.$$

Here conv denotes the convex hull of a set.

Next, we explain how to update the model parameters θ using a subgradient with respect to $w^*(\hat{c})$. If $w^*(\hat{c})$ is differentiable with respect to $\hat{c}$, the chain rule yields

$$\nabla_{\hat{c}} \mathcal{L}_{\text{rob}}(\hat{c}, c_{\text{obs}}) = J_{w^*(\hat{c})}^\top c^*, \quad c^* \in \partial_{w^*(\hat{c})} \mathcal{L}_{\text{rob}}(\hat{c}, c_{\text{obs}}),$$

where $J_{w^*(\hat{c})}$ denotes the Jacobian of $w^*(\hat{c})$ with respect to $\hat{c}$. However, differentiating through the argmin map $w^*(\hat{c})$ is often intractable. Therefore, we employ existing techniques to compute the surrogate of $\nabla_{\hat{c}} \mathcal{L}_{\text{rob}}(\hat{c}, c_{\text{obs}})$ [Elmachtoub and Grigas, 2022, Blondel et al., 2020, Huang and Gupta, 2024]. See Appendix E.2.1 for details of the training algorithm.

3.1.3 Computationally Efficient Training via Approximate Robust Loss

To compute the gradient $\nabla_{\hat{c}} \mathcal{L}_{\text{rob}}(\hat{c}, c_{\text{obs}})$, we use the worst-case coefficient vector $c^* \in \arg \max_{c \in \mathcal{C}} \text{Reg}(\hat{c}, c)$ as an element of the subgradient in practice. However, calculating the worst-case coefficient vector is computationally intensive and incurs significant processing time due to solving a bilevel or nonconvex optimization problem at every training iteration. In this section, we propose a computationally efficient surrogate loss that avoids this costly maximization.

Specifically, given the current optimal solution $w^*(\hat{c})$, we first select a coefficient vector $c^\dagger$ that maximizes the objective value $c^\top w^*(\hat{c})$ over the uncertainty set $\mathcal{C}(c_{\text{obs}})$:

$$c^\dagger(\hat{c}) \in \arg \max_{c \in \mathcal{C}(c_{\text{obs}})} c^\top w^*(\hat{c}). \quad (6)$$

Using $c^\dagger(\hat{c})$, we define the following approximate robust regret loss:

$$\tilde{\mathcal{L}}_{\text{rob}}(\hat{c}, c_{\text{obs}}) := c^\dagger(\hat{c})^\top w^*(\hat{c}) - z(c^\dagger(\hat{c})). \quad (7)$$

As a result, because the optimal value function $z(c)$ no longer appears inside the maximization, the inner problem at each iteration reduces to maximizing a linear function of c for a fixed vector $w^*(\hat{c})$. Moreover, for both ellipsoidal and budgeted uncertainty sets [Ben-Tal et al., 2009, Bertsimas and Sim, 2004], the maximizer $c^\dagger$ admits a closed-form solution; see Appendix B.1.3 for the detailed derivations.

This approximation substantially reduces the computational cost. At each iteration, we only need to solve the optimization problem that produces the decision $w^*(\hat{c})$ and avoid solving the additional nonconvex problem required by $\mathcal{L}_{\text{rob}}$. Therefore, using (7) as the training loss enables robust learning against perturbations that degrade solution quality, while keeping the overall training time manageable. Moreover, gradients can be computed in the same manner as in the previous section, using Danskin’s theorem together with existing techniques [Elmachtoub and Grigas, 2022, Blondel et al., 2020, Huang and Gupta, 2024]. Intuitively, $\tilde{\mathcal{L}}_{\text{rob}}$ (ropt) pushes the model toward perturbation-robust decisions because it trains against the coefficient vector in the uncertainty set that most degrades the current decision $w^*(\hat{c})$. Although it lacks the formal worst-case guarantees of $\mathcal{L}_{\text{rob}}$, it is a computationally efficient alternative.

3.2 DECISION-FOCUSED LEARNING ROBUST TO DISTRIBUTION SHIFTS

In this section, rather than treating the empirical distribution induced by the observed data as the true distribution, we consider an uncertainty set of candidate true distributions around the empirical distribution. We then adopt a DRO framework that minimizes the expected regret under the worst-case distribution in this set. Specifically, we construct a Wasserstein ambiguity set defined using the Wasserstein distance between distributions and derive a DFL loss that is robust to distribution shifts. We first define the conditional empirical distribution based on the dataset $\mathcal{D} := \{(x_i, c_i)\}_{i=1}^N$. We then formulate the loss using the Wasserstein distance.

Here, we denote by δ_c the Dirac measure concentrated at c , i.e., the distribution that assigns probability one to c . Given a covariate value x_i in the dataset, define the index set of samples sharing the same covariate as $\mathcal{I}(x_i) := \{j \in$

$\{1, \dots, N\} \mid \mathbf{x}_j = \mathbf{x}_i\}$. The conditional empirical distribution of the coefficient vector given $\mathbf{x}_i$ is then

$$P_{\tilde{\mathbf{c}}|\mathbf{x}_i} := \frac{1}{|\mathcal{I}(\mathbf{x}_i)|} \sum_{j \in \mathcal{I}(\mathbf{x}_i)} \delta_{\mathbf{c}_j}, \quad (|\mathcal{I}(\mathbf{x}_i)| \geq 1).$$

In many applications, the same covariate value may not appear multiple times in the training data; hence $\mathcal{I}(\mathbf{x}_i) = \{i\}$ in most cases, and $P_{\tilde{\mathbf{c}}|\mathbf{x}_i}$ reduces to the point mass at $\mathbf{c}_i$. To learn a model that is robust to unknown distribution shifts, we consider an ambiguity set that spreads this point-mass distribution and evaluate the worst-case regret over that set.

3.2.1 Robust Loss against Distribution Shift

First, under the conditional empirical distribution of the coefficient vector, the empirical risk can be written as

$$\frac{1}{N} \sum_{i=1}^N \mathbb{E}_{P_{\tilde{\mathbf{c}}|\mathbf{x}_i}} [\text{Reg}(\hat{\mathbf{c}}, \mathbf{c})]. \quad (8)$$

This treats $\text{Reg}(\hat{\mathbf{c}}, \tilde{\mathbf{c}})$ as the per-sample loss, takes the expectation over $\tilde{\mathbf{c}} \sim P_{\tilde{\mathbf{c}}|\mathbf{x}_i}$ for each fixed covariate value $\mathbf{x}_i$, and then averages over $i = 1, \dots, N$. However, at test time, the conditional distribution of the coefficient vector given $\mathbf{x}$ may differ from that observed during training. Minimizing (8) can then overfit to the conditional empirical distribution induced by the training data, potentially increasing regret under the true conditional distribution. To address this issue, we introduce, for each $\mathbf{x}_i$, an uncertainty set of conditional distributions and aim to control the worst-case expected regret over that set.

Specifically, for each covariate value $\mathbf{x}_i$, we measure the discrepancy between conditional distributions of coefficient vectors using the 1-Wasserstein distance:

$$W_1(Q_{\tilde{\mathbf{c}}|\mathbf{x}_i}, P_{\tilde{\mathbf{c}}|\mathbf{x}_i}) := \inf_{\pi \in \Pi(Q_{\tilde{\mathbf{c}}|\mathbf{x}_i}, P_{\tilde{\mathbf{c}}|\mathbf{x}_i})} \int \|\tilde{\mathbf{c}} - \tilde{\mathbf{c}}'\| d\pi(\tilde{\mathbf{c}}, \tilde{\mathbf{c}}').$$

Here, $\Pi(Q_{\tilde{\mathbf{c}}|\mathbf{x}_i}, P_{\tilde{\mathbf{c}}|\mathbf{x}_i})$ denotes the set of couplings π whose marginals are $Q_{\tilde{\mathbf{c}}|\mathbf{x}_i}$, $P_{\tilde{\mathbf{c}}|\mathbf{x}_i}$, and $\|\cdot\|$ denotes a norm. The integral $\int \|\tilde{\mathbf{c}} - \tilde{\mathbf{c}}'\| d\pi(\tilde{\mathbf{c}}, \tilde{\mathbf{c}}')$ is the expected transportation cost under π .

Using this distance, we model uncertainty in the conditional distribution of the coefficient vector via a 1-Wasserstein ambiguity set:

$$\mathbb{B}_\varepsilon(P_{\tilde{\mathbf{c}}|\mathbf{x}_i}) := \left\{ Q_{\tilde{\mathbf{c}}|\mathbf{x}_i} \mid W_1(Q_{\tilde{\mathbf{c}}|\mathbf{x}_i}, P_{\tilde{\mathbf{c}}|\mathbf{x}_i}) \leq \varepsilon \right\}.$$

We then aim to control the worst-case expected regret over this set. For a fixed covariate value $\mathbf{x} = \mathbf{x}_i$, we define the distributionally robust regret loss as

$$\mathcal{L}_{\text{dro}}(\mathbf{x}) := \sup_{Q_{\tilde{\mathbf{c}}|\mathbf{x}} \in \mathbb{B}_\varepsilon(P_{\tilde{\mathbf{c}}|\mathbf{x}})} \mathbb{E}_{\tilde{\mathbf{c}} \sim Q_{\tilde{\mathbf{c}}|\mathbf{x}}} [\text{Reg}(\hat{\mathbf{c}}, \tilde{\mathbf{c}})].$$

Finally, the empirical regret loss minimized is

$$\frac{1}{N} \sum_{i=1}^N \mathcal{L}_{\text{dro}}(\mathbf{x}_i).$$

In what follows, since the empirical regret loss is an average over samples, we fix an arbitrary $\mathbf{x}_i$ and denote it by $\mathbf{x}$ for simplicity. Because $\mathcal{L}_{\text{dro}}(\mathbf{x})$ is an infinite-dimensional optimization problem over probability distributions, directly computing its optimizer is generally intractable. We therefore reduce it into a finite-dimensional optimization problem that is tractable for commercial solvers. Define the regret evaluation function for a fixed prediction $\hat{\mathbf{c}}$ as:

$$\phi_{\mathbf{w}^*(\hat{\mathbf{c}})}(\mathbf{c}) := \text{Reg}(\hat{\mathbf{c}}, \mathbf{c}).$$

Using the strong duality result of Blanchet and Murthy [2019], we show that the distributionally robust regret loss can be written as the sum of the empirical loss and a regularization term, as stated in the following theorem. (See Appendix D.3 for the proof.)

Theorem 3.2. *Let $\mathcal{W}$ be a nonempty compact set. Then the distributionally robust regret loss $\mathcal{L}_{\text{dro}}(\mathbf{x})$ over the 1-Wasserstein ambiguity set $\mathbb{B}_\varepsilon(P_{\tilde{\mathbf{c}}|\mathbf{x}})$ defined with a norm $\|\cdot\|$ admits the following representation:*

$$\mathcal{L}_{\text{dro}}(\mathbf{x}) = \mathbb{E}_{\tilde{\mathbf{c}} \sim P_{\tilde{\mathbf{c}}|\mathbf{x}}} [\phi_{\mathbf{w}^*(\hat{\mathbf{c}})}(\tilde{\mathbf{c}})] + \varepsilon \text{Lip}(\mathbf{w}^*(\hat{\mathbf{c}})). \quad (9)$$

Here, $\text{Lip}(\mathbf{w}^*(\hat{\mathbf{c}}))$ denotes the Lipschitz constant of $\mathbf{w}^*(\hat{\mathbf{c}})$ and is given by $\text{Lip}(\mathbf{w}^*(\hat{\mathbf{c}})) := \sup_{\mathbf{w} \in \mathcal{W}} \|\mathbf{w}^*(\hat{\mathbf{c}}) - \mathbf{w}\|_*$, where $\|\cdot\|_*$ is the dual norm of $\|\cdot\|$.

This theorem provides an intuitive interpretation of the proposed distributionally robust regret. Minimizing the worst-case regret over the Wasserstein ambiguity set is equivalent to minimizing a weighted sum of the expected regret under the empirical distribution induced by the available data and a regularization term that encourages a smaller Lipschitz constant $\text{Lip}(\mathbf{w}^*(\hat{\mathbf{c}}))$. Here, $\text{Lip}(\mathbf{w}^*(\hat{\mathbf{c}}))$ can be interpreted as a sensitivity measure that quantifies how much the regret may change in response to small perturbations of the coefficient vector $\mathbf{c}$. If $\text{Lip}(\mathbf{w}^*(\hat{\mathbf{c}}))$ is large, even a mild distribution shift can lead to a rapid increase in regret. Conversely, keeping this value small promotes stability of regret under distributional changes. Therefore, the proposed method can be viewed as learning a balance between decision quality on the observed data and robustness to distribution shifts.

During training, we evaluate the distributionally robust loss for each sample $\mathbf{x}_i$ and minimize $\frac{1}{N} \sum_{i=1}^N \mathcal{L}_{\text{dro}}(\mathbf{x}_i)$. Using Theorem 3.2, the resulting empirical loss can be written as

$$\frac{1}{N} \sum_{i=1}^N \left\{ \text{Reg}(\hat{\mathbf{c}}(\mathbf{x}_i), \mathbf{c}_i) + \varepsilon \text{Lip}(\mathbf{w}^*(\hat{\mathbf{c}}(\mathbf{x}_i))) \right\}.$$

In other words, the empirical loss decomposes into the standard empirical regret term plus the expected Lipschitz regularizer. As a result, we can incorporate distributional robustness without explicitly solving for a worst-case distribution.

3.2.2 Gradient Calculation of Distributionally Robust Loss

To compute gradients for the proposed distributionally robust regret loss $\mathcal{L}_{\text{dro}}(\mathbf{x})$, we use the decomposition from the previous section

$$\mathcal{L}_{\text{dro}}(\mathbf{x}) = \mathbb{E}_{\tilde{\mathbf{c}} \sim P_{\tilde{\mathbf{c}}|\mathbf{x}}} [\phi_{\mathbf{w}^*(\tilde{\mathbf{c}})}(\tilde{\mathbf{c}})] + \varepsilon \text{Lip}(\mathbf{w}^*(\hat{\mathbf{c}})),$$

and differentiate the two terms separately. The expectation term coincides with the standard DFL objective, so existing surrogate gradients such as SPO+ can be applied directly.

For the second term, $\text{Lip}(\mathbf{w}^*(\hat{\mathbf{c}}))$ can be written as

$$\text{Lip}(\mathbf{w}^*(\hat{\mathbf{c}})) = \max_{(\mathbf{w}, \mathbf{g}) \in \mathcal{W} \times \mathcal{G}} \mathbf{g}^\top (\mathbf{w}^*(\hat{\mathbf{c}}) - \mathbf{w}),$$

where $\mathcal{G}$ denotes the unit ball associated with the dual norm. Hence, its (sub)gradient can be obtained by applying Danskin's theorem. See Appendix C.2 for the detailed derivation and the resulting learning algorithm.

4 THEORETICAL GUARANTEES

We provide generalization guarantees for the proposed robust regret loss and the distributionally robust regret loss. Let $\mathcal{P}$ be the joint distribution over $\mathcal{X} \times \mathbb{R}^n$. In particular, we show that these losses are uniformly bounded and derive a uniform generalization error bound via Rademacher complexity. (see Appendix B.2.)

Theorem 4.1 (Uniform Generalization Error Bound). *Let $\mathcal{F} = \{f_\theta : \theta \in \Theta\}$ be a class of predictive models. Define the loss class $\mathcal{G}_{\text{rob}}$ associated with the robust regret loss as $\mathcal{G}_{\text{rob}} := \{g_\theta : (\mathbf{x}, \mathbf{c}_{\text{obs}}) \mapsto \mathcal{L}_{\text{rob}}(f_\theta(\mathbf{x}), \mathbf{c}_{\text{obs}}) \mid f_\theta \in \mathcal{F}\}$. Under Assumption B.3, for any $\delta \in (0, 1)$, with probability at least $1 - \delta$ on a dataset $\mathcal{D}$, the following generalization bound holds uniformly for all $\theta \in \Theta$:*

$$R(\theta) - \hat{R}_{\mathcal{D}}(\theta) \leq 2\hat{\mathfrak{R}}_{\mathcal{D}}(\mathcal{G}_{\text{rob}}) + 3B\sqrt{\frac{\log(2/\delta)}{2N}}, \quad (10)$$

where $R(\theta)$ and $\hat{R}_{\mathcal{D}}(\theta)$ are the expected risk (with respect to $\mathcal{P}$) and empirical risk defined with respect to the loss $\mathcal{L}_{\text{rob}}$, $\hat{\mathfrak{R}}_{\mathcal{D}}(\mathcal{G}_{\text{rob}})$ is the empirical Rademacher complexity of $\mathcal{G}_{\text{rob}}$, and B is a uniform upper bound on the loss.

An analogous bound holds for the distributionally robust regret loss $\mathcal{L}_{\text{dro}}$ by defining the corresponding loss class $\mathcal{G}_{\text{dro}}$. (see Appendix C.1.)

5 NUMERICAL EXPERIMENTS

In this section, we evaluate the proposed method through experiments on synthetic and real-world datasets. All experiments are fully reproducible, and our implementation is available at <https://github.com/isct-nakatalab/UAI2026-Robust-Decision-Focused-Learning-via-Worst-Case-Regret-Minimization>.

5.1 EXPERIMENTAL SETTINGS

Optimization Problems and Datasets In our experiments, we considered three optimization problems: Shortest Path (SP), Traveling Salesman Problem (TSP) and Knapsack Problem (KP). For the synthetic-data experiments, we used PyEPO [Tang and Khalil, 2024], a DFL library. Specifically, we generated coefficient vectors $\mathbf{c}_i$ via a nonlinear transformation of covariates $\mathbf{x}_i \sim \mathcal{N}(\mathbf{0}, I_p)$. See Appendix F.2 and E.1 for the detailed problem formulations and data-generation procedures. For the real-world experiments, we used the energy scheduling problem (EN) constructed from public datasets [Meinecke et al., 2020, Panadero, 2024, Lamprecht, 2026], following the experimental setup in PyEPO. See Appendix E.1.4 for the detailed formulation of the energy scheduling problem.

Predictive Models and Learning Algorithms For the predictive model f_θ , we used a linear model in the synthetic-data experiments. For the real-world experiments, we used a multilayer perceptron with three fully connected layers (hidden dimensions: 64 and 32) and ReLU activations. We employ the technique from Elmachtoub and Grigas [2022] to compute a surrogate for the gradient of the proposed losses. Detailed hyperparameter settings and the data split are provided in Appendix F.1.

Evaluation Metrics At test time, for each instance i , we compute the optimal solution under the true coefficient vector $\mathbf{w}^*(\mathbf{c}_i)$, and the optimal solution induced by the prediction $\mathbf{w}^*(\hat{\mathbf{c}}_i)$, where $\hat{\mathbf{c}}_i = f_\theta(\mathbf{x}_i)$. Following Section 3.1, we define the regret as $r_i := \text{Reg}(\hat{\mathbf{c}}_i, \mathbf{c}_i)$. To account for scale differences, we report the relative regret $\tilde{r}_i := \frac{r_i}{\|\mathbf{z}(\mathbf{c}_i)\|}$ and use the mean relative regret over the test set as the primary evaluation metric. To assess robustness, we additionally report the Conditional Value-at-Risk (CVaR $_\alpha$) at confidence level $\alpha = 0.9$.

Comparison Methods The methods compared in this experiment are summarized as follows. For clarity, we denote the method names corresponding to our proposed losses below:

- **reg**: The robust regret loss $\mathcal{L}_{\text{rob}}$, minimizing worst-case regret over an uncertainty set.
- **ropt**: The approximate robust regret loss $\tilde{\mathcal{L}}_{\text{rob}}$, a computationally efficient surrogate.
- **droreg**: The distributionally robust regret loss $\mathcal{L}_{\text{dro}}$, minimizing worst-case expected regret over a Wasserstein ambiguity set.

We append `_budget` or `_ellip` to indicate the use of budgeted or ellipsoidal uncertainty sets, respectively. For **droreg**, we specify the ground norm ℓ_p (e.g., **droreg-l2**).

Baselines include:

- **twostage**: A predictive model trained by minimizing mean squared error (MSE).

- **emp**: A predictive model trained with the standard regret loss (3).
- **knn**: A predictive model trained with the local average regret loss based on k -nearest neighbors [Schutte et al., 2024]. We set $k = 5$.
- **gendfl**: A predictive model trained with the regret loss augmented by a generative model [Wang et al., 2026]. Following the original paper, we train the generative model using conditional normalizing flows [Winkler et al., 2019] and set $\alpha = 0.9$ for CVaR.

5.1.1 Uncertainty and Perturbation Settings

In this experiment, we use the following perturbation settings to mimic two types of uncertainty: observation errors and distribution shifts.

Observation Error Settings To artificially introduce observation errors into the training data, for each training instance i we generate a perturbed observed coefficient vector $\mathbf{c}_{\text{obs},i}$ from the true coefficient vector $\mathbf{c}_i$ and use it for training. We consider the following two perturbation types.

- **Ellipsoidal Perturbation (ell)**: For each instance i , we set the baseline standard deviation as $\sigma_i^{(\text{base})} := \|\mathbf{c}_i\|_2 / \sqrt{n}$. We then draw Gaussian noise $\mathbf{g}_i \sim \mathcal{N}(\mathbf{0}, (\sigma_i^{(\text{base})})^2 I_n)$ and define the observed coefficient vector as $\mathbf{c}_{\text{obs},i} = \mathbf{c}_i + \mathbf{g}_i$.
- **Sparse Perturbation (sparse)**: Sparse perturbations selectively introduce large shifts in a subset of coordinates. Specifically, we add a fixed shift of width $\delta \geq 0$ to the randomly selected $\lfloor \gamma_{\text{frac}} n \rfloor$ coordinates, where $\gamma_{\text{frac}} \in (0, 1]$ controls the fraction of perturbed dimensions. See Appendix F.1 for details. For each instance i , let $J_i \subset \{1, \dots, n\}$ be the indices of the selected coordinates and $s \in \{+1, -1\}$ denote the problem direction, where $s = +1$ for the maximization problems and $s = -1$ for the minimization problems. Then we set $c_{\text{obs},ij} = c_{ij} + s \delta \mathbf{1}\{j \in J_i\}$.

Distribution Shift Settings To evaluate robustness to distribution shifts, we artificially alter the conditional distribution of coefficient vectors in the training data. Concretely, we use a mixture perturbation model that combines two types of shifts: covariate-dependent bias, where a polynomial map of the covariates is added to the coefficients to simulate systematic shifts and global noise inflation, where covariate-independent Gaussian noise is injected to simulate increased environmental variance. Detailed mathematical formulations and parameter values are provided in Appendix F.1.

5.2 EXPERIMENTAL RESULTS AND DISCUSSION

5.2.1 Robustness against Observation Errors

Figure 1 (a) reports the mean relative regret for each optimization problem and perturbation setting on the test set.

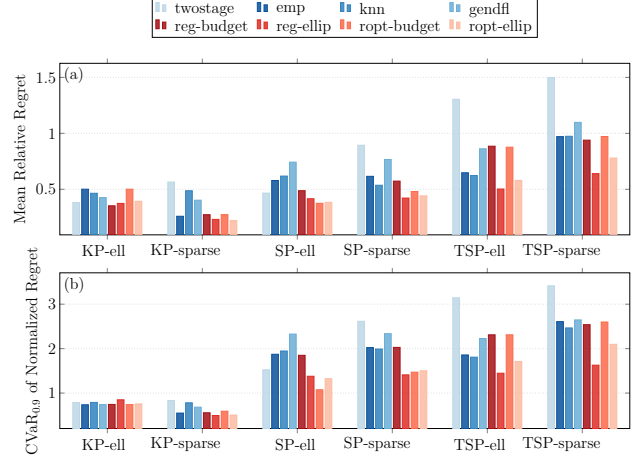


Figure 1: Mean relative regret (a) and CVaR_{0.9} (b) in the presence of observation errors.

Since **twostage** directly minimizes MSE, it does not show consistent improvements over **emp**, which is trained with a regret-based loss; in particular, it performs substantially worse on TSP. We also observe that **gendfl** yields a relatively large mean relative regret overall. In the presence of observation errors in the training data, the generative model learns a distorted conditional distribution, causing tail-focused objectives (e.g., CVaR) to target the wrong tail and leading to evaluation under a misspecified distribution. In contrast, the proposed methods (**reg** and **ropt**) achieve lower regret than **emp** and **knn** in many settings. Notably, **reg-ellip** provides a marked improvement on TSP, highlighting the benefit of optimizing a loss that accounts for worst-case outcomes. Moreover, even under mismatched uncertainty settings, where the ellipsoidal-set models (**reg-ellip/ropt-ellip**) are evaluated under structurally different sparse perturbations, they still tend to outperform **emp** and **knn**. This suggests that a model trained with the ellipsoidal set can remain robust even when the perturbations encountered at evaluation differ from those assumed during training.

Figure 1 (b) shows the resulting CVaR_{0.9} values. Beyond improving the average regret, the proposed methods reduce CVaR_{0.9} by suppressing a small number of extreme regret instances. This reduction in tail risk also contributes to the lower mean relative regret. Conversely, methods without explicit worst-case robustness, such as **knn**, exhibit higher CVaR_{0.9}, which is consistent with their increased mean relative regret.

5.2.2 Robustness against Distribution Shifts

Figure 2 compares the mean relative regret and CVaR_{0.9} for each optimization problem under distribution shift. Overall, **droreg** exhibits more stable performance than the baselines in terms of both mean relative regret and CVaR_{0.9}. This is consistent with the fact that **twostage**, which minimizes only prediction error, is sensitive to distributional

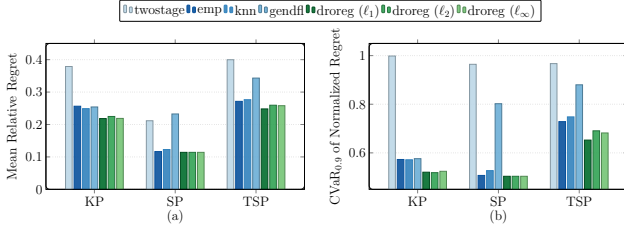


Figure 2: Comparison of mean relative regret (a) and $\text{CVaR}_{0.9}$ (b) for each optimization problem under distribution shift.

changes, while `emp` can degrade out-of-distribution due to overfitting to the training distribution. Similarly, `gendfl` deteriorates when the generative model fails to capture the shifted distribution. In contrast, the proposed distributionally robust approach yields substantial gains on KP and TSP; notably, using the ℓ_1 leads to particularly strong robustness on TSP. We note that the synthetic shifts were generated by a specific mechanism that was not assumed during training (Appendix F.1). Even so, `droreg` still outperformed the baselines, indicating that it remains robust even when the deployment shifts are not the ones assumed during training.

5.2.3 Sensitivity Analysis

We show the results of sensitivity analyses on the uncertainty-set hyperparameters, using `reg-ellip` on TSP under sparse perturbation as a representative case (Figure 3). Both the mean and $\text{CVaR}_{0.9}$ of the relative regret are lowest at an intermediate κ and larger when κ is too small or too large. Across most of the range, the proposed method outperforms the `emp` and `knn` baselines in both metrics, matching them only at the smallest κ . We also show the results for `droreg` by varying the parameter ε for the Wasserstein ambiguity set with several ground norms (ℓ_1 , ℓ_2 , ℓ_∞) on KP (Figure 4). Both the mean and $\text{CVaR}_{0.9}$ of the relative regret are lowest at an intermediate ε and larger when ε is too small or too large, especially under the ℓ_2 and ℓ_∞ norms, whereas the ℓ_1 norm is the most stable across ε . With an appropriately chosen ε , `droreg` outperforms the baselines, and the best ground norm depends on the problem structure. Comprehensive results of the sensitivity analysis are shown in Appendix G (Figure 9).

5.2.4 Comparison of Computation Time

We investigate how much faster `ropt` is compared to other methods to verify its computational efficiency. Specifically, we compare the computation times of the existing methods (`emp`, `knn`) and the proposed methods (`ropt`, `reg`) for robust regret learning against observation errors. Figure 5 shows the distribution of per-epoch training time (the vertical axis is on a logarithmic scale). Both `emp` and `knn` have low computational cost because they only require solving the downstream optimization problem for the predicted coefficient vector. In contrast, `reg` computes the worst-case

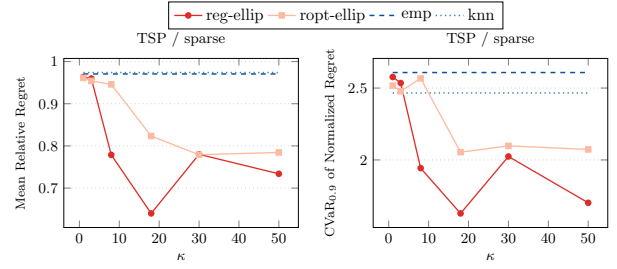


Figure 3: Sensitivity of mean relative regret (left) and $\text{CVaR}_{0.9}$ (right) for `reg-ellip` (`ropt-ellip` for reference) to the parameter κ on TSP under sparse perturbation, with the `emp` (dashed) and `knn` (dotted) baselines.

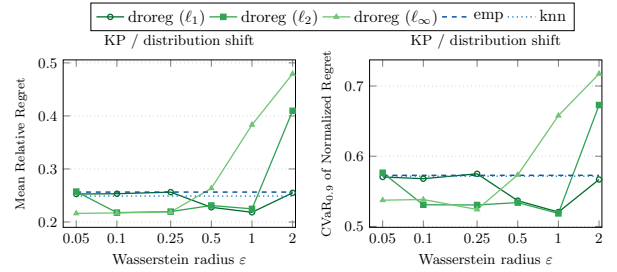


Figure 4: Sensitivity of mean relative regret (left) and $\text{CVaR}_{0.9}$ (right) for `droreg` with ℓ_1 , ℓ_2 , ℓ_∞ ground norms to the parameter ε on KP under distribution shift, with the `emp` (dashed) and `knn` (dotted) baselines.

regret exactly, incurring the longest training time; `ropt`, in turn, takes more time than `emp/knn` but far less than `reg`, while still improving robustness. Overall, `ropt` provides a practical trade-off, enhancing robustness within feasible computation time.

5.2.5 Verification using Real-World Data

The results on the real-world energy scheduling problem (EN) are shown in Figure 6 (observation errors) and Figure 7 (distribution shift).

Under Observation Error In the presence of observation errors, the proposed methods consistently outperform the baselines in both mean relative regret and $\text{CVaR}_{0.9}$. This indicates that the benefits of our loss functions observed on synthetic data also carry over to a more realistic setting based on public datasets. Notably, the improvement in $\text{CVaR}_{0.9}$ suggests that the proposed methods not only reduce regret but also suppress rare yet severe regret events.

Under Distribution Shift Even under distribution shift, the proposed methods tend to keep both mean relative regret and $\text{CVaR}_{0.9}$ low. In practice, electricity prices and demand can shift due to seasonality, system and policy changes, or evolving market conditions, making risk reduction under distribution shift particularly important. Overall, real-world experiments demonstrate that the proposed methods improve solution robustness in terms of both regret and CVaR

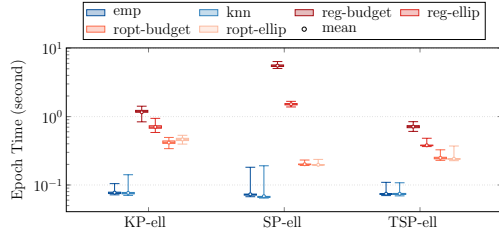


Figure 5: Comparison of computation time using the proposed robust regret.

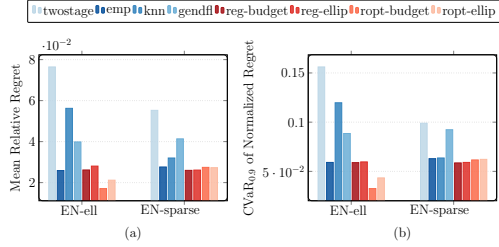


Figure 6: Comparison of mean relative regret (a) and $CVaR_{0.9}$ (b) under observation errors in real-world data.

against two distinct sources of uncertainty.

6 CONCLUSION

We proposed robust DFL methods that address (i) observation errors in training coefficient vectors and (ii) distribution shifts at deployment. For Uncertainty (i), we define an uncertainty set for observation errors and minimize the worst-case regret over the set. For Uncertainty (ii), we adopt a Wasserstein ambiguity set and optimize a distributionally robust regret. Experiments show that our method improves regret and risk measures over existing baselines. Future work includes automating the selection of uncertainty-set hyperparameters.

Limitations and Future Work Several limitations remain. First, using the approximate loss $\tilde{\mathcal{L}}_{\text{rob}}$ (ropt) in place of the exact loss $\mathcal{L}_{\text{rob}}$ is a heuristic without formal worst-case guarantees. Although it performed comparably to the exact loss at much lower computation cost, establishing such guarantees remains an open problem. Second, our real-world validation relies on a single energy-scheduling dataset, and our synthetic shifts may not fully capture the shifts arising in practice, so broader evaluation is needed. Third, automating the data-driven selection of the parameters of the uncertainty and ambiguity sets is an important direction for future work.

References

Akshay Agrawal, Brandon Amos, Shane Barratt, Stephen Boyd, Steven Diamond, and J. Zico Kolter. Differen-

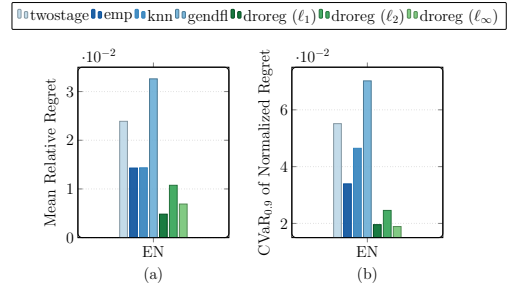


Figure 7: Comparison of mean relative regret (a) and $CVaR_{0.9}$ (b) under distribution shifts in real-world data.

tiatile convex optimization layers. In *Advances in Neural Information Processing Systems*, volume 32, 2019.

Brandon Amos and J. Zico Kolter. OptNet: Differentiable optimization as a layer in neural networks. In *Proceedings of the 34th International Conference on Machine Learning*, volume 70 of *Proceedings of Machine Learning Research*, pages 136–145, 2017.

Peter L. Bartlett and Shahar Mendelson. Rademacher and gaussian complexities: Risk bounds and structural results. *Journal of Machine Learning Research*, 3:463–482, 2002.

Aharon Ben-Tal, Laurent El Ghaoui, and Arkadi Nemirovski. *Robust Optimization*. Princeton Series in Applied Mathematics. Princeton University Press, Princeton, NJ, 2009. doi: 10.1515/9781400831050.

Yoshua Bengio, Andrea Lodi, and Antoine Prouvost. Machine learning for combinatorial optimization: A methodological tour d’horizon. *European Journal of Operational Research*, 290(2):405–421, 2021. doi: 10.1016/j.ejor.2020.07.063.

Dimitri P. Bertsekas. *Nonlinear Programming*. Athena Scientific, Belmont, MA, 2 edition, 1999. ISBN 1886529000.

Dimitris Bertsimas and Nathan Kallus. From predictive to prescriptive analytics. *Management Science*, 66(3): 1025–1044, 2020. doi: 10.1287/mnsc.2018.3253.

Dimitris Bertsimas and Melvyn Sim. The price of robustness. *Operations Research*, 52(1):35–53, 2004. doi: 10.1287/opre.1030.0065.

Dimitris Bertsimas and John N. Tsitsiklis. *Introduction to Linear Optimization*. Athena Scientific, 1997.

Dimitris Bertsimas, David B. Brown, and Constantine Caramanis. Theory and applications of robust optimization. *SIAM Review*, 53(3):464–501, 2011. doi: 10.1137/080734510.

Jose Blanchet and Karthyek Murthy. Quantifying distributional model risk via optimal transport. *Mathematics of Operations Research*, 44(2):565–600, 2019. doi: 10.1287/moor.2018.0936.

- Mathieu Blondel, André F. T. Martins, and Vlad Niculae. Learning with fenchel-young losses. *Journal of Machine Learning Research*, 21(35):1–69, 2020. URL <https://jmlr.org/papers/v21/19-021.html>.
- Erick Delage and Yinyu Ye. Distributionally robust optimization under moment uncertainty with application to data-driven problems. *Operations Research*, 58(3): 595–612, 2010. doi: 10.1287/opre.1090.0741.
- Priya L. Donti, Brandon Amos, and J. Zico Kolter. Task-based end-to-end model learning in stochastic optimization. In *Proceedings of the 31st International Conference on Neural Information Processing Systems, NIPS’17*, page 5490–5500, Red Hook, NY, USA, 2017. Curran Associates Inc. ISBN 9781510860964.
- Othman El Balghiti, Adam N. Elmachtoub, Paul Grigas, and Ambuj Tewari. Generalization bounds in the predict-then-optimize framework. *Mathematics of Operations Research*, 48(4):2043–2065, 2023. doi: 10.1287/moor.2022.1330. URL <https://doi.org/10.1287/moor.2022.1330>.
- Adam N. Elmachtoub and Paul Grigas. Smart “predict, then optimize”. *Management Science*, 68(1):9–26, 2022. doi: 10.1287/mnsc.2020.3922. URL <https://doi.org/10.1287/mnsc.2020.3922>.
- Aaron M. Ferber, Bryan Wilder, Bistra Dilkina, and Milind Tambe. MIPaaL: Mixed integer program as a layer. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 34, pages 1504–1511, 2020. doi: 10.1609/aaai.v34i02.5509.
- Gurobi Optimization, LLC. Gurobi optimizer reference manual. Software documentation, 2025. URL <https://docs.gurobi.com/projects/optimizer/en/current/>. (2026-01-21).
- Frederick S. Hillier and Gerald J. Lieberman. *Introduction to Operations Research*. McGraw-Hill Education, 10 edition, 2015.
- Michael Huang and Vishal Gupta. Decision-focused learning with directional gradients. In *Advances in Neural Information Processing Systems*, volume 37, 2024.
- Jakob Huber, Sebastian Müller, Moritz Fleischmann, and Heiner Stuckenschmidt. A data-driven newsvendor problem: From data to decision. *European Journal of Operational Research*, 278(3):904–915, 2019. doi: 10.1016/j.ejor.2019.04.043.
- Diederik P. Kingma and Jimmy Ba. Adam: A method for stochastic optimization. In *International Conference on Learning Representations (ICLR)*, 2015. URL <https://arxiv.org/abs/1412.6980>.
- James Kotary, Ferdinando Fioretto, Pascal Van Hentenryck, and Bryan Wilder. End-to-end constrained optimization learning: A survey. In *Proceedings of the 30th International Joint Conference on Artificial Intelligence*, pages 4475–4482, 2021. doi: 10.24963/ijcai.2021/610.
- Christian Sebastian Lamprecht. meteostat: Access and analyze historical weather and climate data with python. Software (Python package / PyPI), 2026. URL <https://pypi.org/project/meteostat/>. (2026-01-21).
- Xutao Ma, Chao Ning, and Wenli Du. Differentiable distributionally robust optimization layers. In *Proceedings of the 41st International Conference on Machine Learning*, 2024.
- Jayanta Mandi, James Kotary, Senne Berden, Maxime Mulamba, Victor Bucarey, Tias Guns, and Ferdinando Fioretto. Decision-focused learning: Foundations, state of the art, benchmark and future opportunities. *Journal of Artificial Intelligence Research*, 2024. doi: 10.1613/jair.1.15320.
- Colin McDiarmid. On the method of bounded differences. *Surveys in Combinatorics*, 141:148–188, 1989.
- Steffen Meinecke, Džanan Sarajlić, Simon Ruben Drauz, Annika Klettke, Lars-Peter Lauven, Christian Rehtanz, Albert Moser, and Martin Braun. Simbench—a benchmark dataset of electric power systems to compare innovative solutions based on power flow analysis. *Energies*, 13(12):3290, 2020. doi: 10.3390/en13123290.
- Arthur Mensch and Mathieu Blondel. Differentiable dynamic programming for structured prediction and attention. In *Proceedings of the 35th International Conference on Machine Learning*, volume 80 of *Proceedings of Machine Learning Research*, pages 3462–3471, 2018.
- 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. doi: 10.1007/s10107-017-1172-1.
- Eugenio Panadero. aiopvpc: Simple aio library to download spanish electricity hourly prices. GitHub repository, 2024. URL <https://github.com/azogue/aiopvpc>. (2026-01-21).
- Ignacio Rios, Roger J.-B. Wets, and David L. Woodruff. Multi-period forecasting and scenario generation with limited data. *Computational Management Science*, 12(2): 267–295, 2015. doi: 10.1007/s10287-015-0230-5.
- Utsav Sadana, Abhilash Chenreddy, Erick Delage, Alexandre Forel, Emma Frejinger, and Thibaut Vidal. A survey of contextual optimization methods for decision-making

under uncertainty. *European Journal of Operational Research*, 320(2):271–289, 2025. doi: 10.1016/j.ejor.2024.03.020.

Noah J. Schutte, Krzysztof S. Postek, and Neil Yorke-Smith. Robust losses for decision-focused learning. In *Proceedings of the 33rd International Joint Conference on Artificial Intelligence*, pages 4868–4875, 2024. doi: 10.24963/ijcai.2024/538.

Sanket Shah, Kai Wang, Bryan Wilder, Andrew Perrault, and Milind Tambe. Decision-focused learning without decision-making: Learning locally optimized decision losses. In *Advances in Neural Information Processing Systems*, volume 35, 2022. URL <https://dblp.org/rec/conf/nips/Shah0WPT22>.

Bo Tang and Elias B. Khalil. PyEPO: a PyTorch-based end-to-end predict-then-optimize library for linear and integer programming. *Mathematical Programming Computation*, July 2024. ISSN 1867-2957. doi: 10.1007/s12532-024-00255-x.

Marin Vlastelica, Anselm Paulus, Vít Musil, Georg Martius, and Michal Rolínek. Differentiation of blackbox combinatorial solvers. In *International Conference on Learning Representations*, 2020.

Prince Zizhuang Wang, Shuyi Chen, Jinhao Liang, Ferdinando Fioretto, and Shixiang Zhu. Gen-DFL: Decision-focused generative learning for robust decision making. In *The Fourteenth International Conference on Learning Representations*, 2026. URL <https://openreview.net/forum?id=GU2197a3Lm>.

Wolfram Wiesemann, Daniel Kuhn, and Melvyn Sim. Distributionally robust convex optimization. *Operations Research*, 62(6):1358–1376, 2014. doi: 10.1287/opre.2014.1314.

Bryan Wilder, Bistra Dilkina, and Milind Tambe. Melding the data-decisions pipeline: Decision-focused learning for combinatorial optimization. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 33, pages 1658–1665, 2019. doi: 10.1609/aaai.v33i01.33011658.

Christina Winkler, Daniel Worrall, Emiel Hoogeboom, and Max Welling. Learning likelihoods with conditional normalizing flows. *arXiv preprint*, 2019. URL <https://arxiv.org/abs/1912.00042>.

Robust Decision-Focused Learning via Worst-Case Regret Minimization (Supplementary Material)

Shoki Yamao¹ Ken Kobayashi¹ Ryo Matsui² Shota Nagai¹ Naoki Nishimura² Kazuhide Nakata¹

¹Institute of Science Tokyo, Tokyo, Japan

²Recruit Co., Ltd., Chiyoda, Tokyo, Japan,

This Supplementary Material should be submitted together with the main paper.

A MATHEMATICAL PRELIMINARIES

In this section, we summarize external results used in this paper to clarify the theoretical background.

Proposition A.1 (Danskin’s Theorem; [Bertsekas, 1999, Proposition B.25]). *Let Y be a non-empty compact set, and let the function $\psi : \mathbb{R}^m \times Y \rightarrow \mathbb{R}$ be continuous, and for each $y \in Y$, $\psi(\cdot, y)$ be convex and of class C^1 . Define the optimal value function $f(x) := \max_{y \in Y} \psi(x, y)$ and the set of optimal solutions $Y(x) := \arg \max_{y \in Y} \psi(x, y)$. In this case, if $\nabla_x \psi(x, \cdot)$ is continuous on Y , the subgradient of $f(x)$ is given by:*

$$\partial f(x) = \text{conv}\{\nabla_x \psi(x, y) \mid y \in Y(x)\}.$$

Here, conv denotes the convex hull. In particular, if the optimal solution is unique, the gradient is $\nabla f(x) = \nabla_x \psi(x, y^*)$.

Theorem A.2 (Blanchet–Murthy’s Strong Duality [Blanchet and Murthy, 2019, Theorem 1]). *Consider a distance $d(\mathbf{u}, \mathbf{v}) := \|\mathbf{u} - \mathbf{v}\|$ induced by a norm $\|\cdot\|$.*

Let μ be a probability measure on $\mathbb{R}^n$, and write $\mathbf{C} \sim \mu$. Assume that the evaluation function $f : \mathbb{R}^n \rightarrow \mathbb{R}$ is

- *absolutely integrable with respect to μ ,*
- *upper semicontinuous on $\mathbb{R}^n$.*

In this case, for any $\varepsilon > 0$, the worst-case expected value on the uncertainty set with respect to W_1

$$\mathbb{B}_\varepsilon(\mu) := \{Q \mid W_1(Q, \mu) \leq \varepsilon\}$$

is expressed as the following univariate dual problem:

$$\sup_{Q \in \mathbb{B}_\varepsilon(\mu)} \mathbb{E}_Q[f(\mathbf{C})] = \inf_{\lambda \geq 0} \left\{ \lambda \varepsilon + \mathbb{E}_\mu \left[\sup_{\mathbf{y} \in \mathbb{R}^n} \{f(\mathbf{y}) - \lambda \|\mathbf{y} - \mathbf{C}\|\} \right] \right\}. \quad (11)$$

The minimum value on the right-hand side is finite and equal to the maximum value on the left-hand side.

B DETAILS OF ROBUST REGRET LOSS

B.1 DETAILED FORMULATION OF ROBUST REGRET LOSS

In this section, we describe the detailed formulation of the robust regret loss $\mathcal{L}_{\text{rob}}$.

As concrete shapes of the uncertainty set $\mathcal{C}(\mathbf{c}_{\text{obs}})$, we adopt the following two sets used in robust optimization.

1. Ellipsoidal Uncertainty Set [Ben-Tal et al., 2009]: It is a set defined based on the Mahalanobis distance, considering the correlation between coefficient vectors.

$$\mathcal{C}_{\text{ell}}(\mathbf{c}_{\text{obs}}) = \{ \mathbf{c} \in \mathbb{R}^n \mid (\mathbf{c} - \mathbf{c}_{\text{obs}})^\top \Sigma^{-1} (\mathbf{c} - \mathbf{c}_{\text{obs}}) \leq \kappa^2 \}.$$

Here, Σ is the covariance matrix of the error, and κ is a parameter that controls the magnitude of uncertainty.

2. Budgeted Uncertainty Set [Bertsimas and Sim, 2004]: It is an uncertainty set where each component can fluctuate independently, but the total amount (budget) of components having large errors at the same time is limited. Specifically,

$$\mathcal{C}_\Gamma(\mathbf{c}_{\text{obs}}) = \{ \mathbf{c} \in \mathbb{R}^n \mid \mathbf{c} = \mathbf{c}_{\text{obs}} + \text{Diag}(\boldsymbol{\rho})\boldsymbol{\zeta}, \|\boldsymbol{\zeta}\|_\infty \leq 1, \|\boldsymbol{\zeta}\|_1 \leq \Gamma \}.$$

Here, $\boldsymbol{\rho} = (\rho_1, \dots, \rho_n)^\top$ is a vector in which the fluctuation range of each component is arranged, and $\text{Diag}(\boldsymbol{\rho}) \in \mathbb{R}^{n \times n}$ represents a diagonal matrix with diagonal elements $\rho_1, \dots, \rho_n$ and 0 otherwise. Also, Γ represents the budget of uncertainty.

First, defining the support function on the uncertainty set as $h_{\mathcal{C}}(\mathbf{d}) := \sup_{\mathbf{c} \in \mathcal{C}} \mathbf{c}^\top \mathbf{d}$, the robust regret loss can be transformed as follows:

$$\mathcal{L}_{\text{rob}}(\hat{\mathbf{c}}, \mathbf{c}_{\text{obs}}) = \max_{\mathbf{c} \in \mathcal{C}(\mathbf{c}_{\text{obs}})} \left\{ \mathbf{c}^\top \mathbf{w}^*(\hat{\mathbf{c}}) - \min_{\mathbf{w} \in \mathcal{W}} \mathbf{c}^\top \mathbf{w} \right\} \quad (12)$$

$$= \max_{\mathbf{c} \in \mathcal{C}(\mathbf{c}_{\text{obs}})} \max_{\mathbf{w} \in \mathcal{W}} \mathbf{c}^\top (\mathbf{w}^*(\hat{\mathbf{c}}) - \mathbf{w}) \quad (13)$$

$$= \max_{\mathbf{w} \in \mathcal{W}} \max_{\mathbf{c} \in \mathcal{C}(\mathbf{c}_{\text{obs}})} \mathbf{c}^\top (\mathbf{w}^*(\hat{\mathbf{c}}) - \mathbf{w}) \quad (14)$$

$$= \max_{\mathbf{w} \in \mathcal{W}} h_{\mathcal{C}(\mathbf{c}_{\text{obs}})}(\mathbf{w}^*(\hat{\mathbf{c}}) - \mathbf{w}). \quad (15)$$

In the following, we show that the solution of the inner optimization by the uncertainty set can be analytically calculated when using the ellipsoidal uncertainty set and the budgeted uncertainty set for Equation (15), and formulate it as an optimization problem.

B.1.1 Formulation with Ellipsoidal Uncertainty Set

We describe the details when using the ellipsoidal uncertainty set $\mathcal{C}_{\text{ell}}(\mathbf{c}_{\text{obs}})$. The support function on this set (the value maximizing the linear function $\mathbf{c}^\top \mathbf{d}$ on the set) can be analytically obtained as follows [Ben-Tal et al., 2009].

$$h_{\mathcal{C}_{\text{ell}}(\mathbf{c}_{\text{obs}})}(\mathbf{d}) = \sup_{\mathbf{c} \in \mathcal{C}_{\text{ell}}(\mathbf{c}_{\text{obs}})} \mathbf{c}^\top \mathbf{d} = \mathbf{c}_{\text{obs}}^\top \mathbf{d} + \kappa \|\Sigma^{1/2} \mathbf{d}\|_2.$$

By substituting this into Equation (15), the calculation of the worst regret becomes the following maximization problem:

$$\mathcal{L}_{\text{rob}} = \max_{\mathbf{w} \in \mathcal{W}} \left\{ \mathbf{c}_{\text{obs}}^\top (\mathbf{w}^*(\hat{\mathbf{c}}) - \mathbf{w}) + \kappa \|\Sigma^{1/2} (\mathbf{w}^*(\hat{\mathbf{c}}) - \mathbf{w})\|_2 \right\}. \quad (16)$$

Furthermore, by introducing an auxiliary variable $t \geq 0$, it can be transformed into the following formally equivalent problem:

$$\begin{aligned} \max_{\mathbf{w} \in \mathcal{W}, t \geq 0} \quad & \mathbf{c}_{\text{obs}}^\top (\mathbf{w}^*(\hat{\mathbf{c}}) - \mathbf{w}) + \kappa t \\ \text{s.t.} \quad & t^2 \leq (\mathbf{w}^*(\hat{\mathbf{c}}) - \mathbf{w})^\top \Sigma (\mathbf{w}^*(\hat{\mathbf{c}}) - \mathbf{w}). \end{aligned} \quad (17)$$

Although this problem is a non-convex Quadratically Constrained Quadratic Programming (QCQP), it is possible to find a globally optimal solution using mathematical optimization solvers such as Gurobi Optimizer [Gurobi Optimization, LLC, 2025].

B.1.2 Formulation with Budgeted Uncertainty Set

Next, we describe the details when using the budgeted uncertainty set $\mathcal{C}_\Gamma(\mathbf{c}_{\text{obs}})$. The support function in this case can be formulated as follows.

$$h_{\mathcal{C}_\Gamma(\mathbf{c}_{\text{obs}})}(\mathbf{d}) = \sup_{\mathbf{c} \in \mathcal{C}_\Gamma(\mathbf{c}_{\text{obs}})} \mathbf{c}^\top \mathbf{d} = \mathbf{c}_{\text{obs}}^\top \mathbf{d} + \max_{\substack{|\zeta_j| \leq 1 \\ \sum_{j=1}^n |\zeta_j| \leq \Gamma}} \sum_{j=1}^n \rho_j \zeta_j d_j.$$

Here, if we transform the variables as $\zeta_j d_j = \text{sgn}(d_j) s_j$ (where $s_j \in [0, 1]$), the maximization problem of the second term becomes

$$\max_{\substack{0 \leq s_j \leq 1 \\ \sum_{j=1}^n s_j \leq \Gamma}} \sum_{j=1}^n \rho_j |d_j| s_j \quad (18)$$

and the optimal solution is obtained by selecting s_j in descending order of weight $\rho_j |d_j|$. Specifically, the support function is expressed as

$$h_{\mathcal{C}_\Gamma(\mathbf{c}_{\text{obs}})}(\mathbf{d}) = \mathbf{c}_{\text{obs}}^\top \mathbf{d} + \sum_{k=1}^{\lfloor \Gamma \rfloor} \rho_{(k)} |d_{(k)}| + (\Gamma - \lfloor \Gamma \rfloor) \rho_{(\lfloor \Gamma \rfloor + 1)} |d_{(\lfloor \Gamma \rfloor + 1)}|.$$

Here, $\rho_{(k)} |d_{(k)}|$ represents the k -th value when $\{\rho_j |d_j|\}_{j=1}^n$ is sorted in descending order. In particular, when Γ is an integer, it simplifies to

$$h_{\mathcal{C}_\Gamma(\mathbf{c}_{\text{obs}})}(\mathbf{d}) = \mathbf{c}_{\text{obs}}^\top \mathbf{d} + \sum_{k=1}^{\Gamma} \rho_{(k)} |d_{(k)}|.$$

By substituting $\mathbf{d} = \mathbf{w}^*(\hat{\mathbf{c}}) - \mathbf{w}$, the robust loss function can be written as

$$\mathcal{L}_{\text{rob}}(\hat{\mathbf{c}}, \mathbf{c}_{\text{obs}}) = \max_{\mathbf{w} \in \mathcal{W}} \left\{ \mathbf{c}_{\text{obs}}^\top (\mathbf{w}^*(\hat{\mathbf{c}}) - \mathbf{w}) + \sum_{k=1}^{\Gamma} \rho_{(k)} |w^*(\hat{\mathbf{c}})_k - w_k| \right\}.$$

Furthermore, when $\mathcal{W}$ is a polyhedral set and Γ is an integer, this maximization problem can be formulated as a Mixed Integer Linear Programming (MILP). We introduce a variable $u_j \geq |w^*(\hat{\mathbf{c}})_j - w_j|$ representing the absolute value of each component, and a 0-1 variable $y_j \in \{0, 1\}$ representing whether it is included in the top Γ . Furthermore, by using a variable t_j to linearize the product term, it can be reduced to the following MILP as a whole:

$$u_j \geq w^*(\hat{\mathbf{c}})_j - w_j, \quad u_j \geq -(w^*(\hat{\mathbf{c}})_j - w_j), \quad u_j \geq 0 \quad (j = 1, \dots, n).$$

We introduce a 0-1 variable $y_j \in \{0, 1\}$ to represent the top Γ selection, and set $\sum_{j=1}^n y_j \leq \Gamma$. To linearize the product $\rho_j u_j y_j$ appearing in the objective function, we introduce a new variable $t_j = \rho_j u_j y_j$. Summarizing the above, the outer maximization reduces to the following MILP:

$$\max_{\substack{\mathbf{w} \in \mathcal{W}, \mathbf{u} \geq 0, \\ \mathbf{y} \in \{0, 1\}^n, \mathbf{t} \geq 0}} \mathbf{c}_{\text{obs}}^\top (\mathbf{w}^*(\hat{\mathbf{c}}) - \mathbf{w}) + \sum_{j=1}^n t_j \quad (19)$$

$$\text{s.t. } u_j \geq w^*(\hat{\mathbf{c}})_j - w_j, \quad u_j \geq -(w^*(\hat{\mathbf{c}})_j - w_j) \quad (j = 1, \dots, n), \quad (20)$$

$$0 \leq u_j \leq M_j \quad (j = 1, \dots, n), \quad (21)$$

$$\sum_{j=1}^n y_j \leq \Gamma, \quad y_j \in \{0, 1\} \quad (j = 1, \dots, n), \quad (22)$$

$$0 \leq t_j \leq \rho_j u_j, \quad t_j \leq \rho_j M_j y_j, \quad t_j \geq \rho_j u_j - \rho_j M_j (1 - y_j) \quad (j = 1, \dots, n). \quad (23)$$

The objective function is $\mathbf{c}_{\text{obs}}^\top (\mathbf{w}^*(\hat{\mathbf{c}}) - \mathbf{w}) + \sum_{j=1}^n t_j$, and since t_j matches $\rho_j u_j$ for the selected Γ components, it is equivalent to the sum of the top Γ of the support function. Note that M_j is a sufficiently large constant satisfying $M_j \geq \max_{\mathbf{w}, \mathbf{w}' \in \mathcal{W}} |w_j - w'_j|$.

B.1.3 Analytical Solution for Approximate Robust Regret Loss

We describe the solution method for the maximization problem over the uncertainty set of the coefficient vector, which is important in the calculation of the approximate robust regret loss $\hat{\mathcal{L}}_{\text{rob}}$:

$$\mathbf{c}^\dagger(\hat{\mathbf{c}}) \in \arg \max_{\mathbf{c} \in \mathcal{C}(\mathbf{c}_{\text{obs}})} \mathbf{c}^\top \mathbf{w}^*(\hat{\mathbf{c}}).$$

This $\mathbf{c}^\dagger$ is used to approximate the inner minimization in the regret calculation.

Also, for any set $\mathcal{C}(\mathbf{c}_{\text{obs}})$, let $\max_{\mathbf{c} \in \mathcal{C}(\mathbf{c}_{\text{obs}})} \mathbf{c}^\top \mathbf{w}^*(\hat{\mathbf{c}}) = h_{\mathcal{C}(\mathbf{c}_{\text{obs}})}(\mathbf{w}^*(\hat{\mathbf{c}}))$.

Case of Budget Uncertainty Set The support function of the budget uncertainty set $\mathcal{C}_\Gamma(\mathbf{c}_{\text{obs}})$ can be written as follows, by a similar argument to Equation (18):

$$h_{\mathcal{C}_\Gamma(\mathbf{c}_{\text{obs}})}(\mathbf{w}) = \mathbf{c}_{\text{obs}}^\top \mathbf{w} + \sum_{k=1}^{\lfloor \Gamma \rfloor} \rho_{(k)} |w_{(k)}| + (\Gamma - \lfloor \Gamma \rfloor) \rho_{(\lfloor \Gamma \rfloor + 1)} |w_{(\lfloor \Gamma \rfloor + 1)}|. \quad (24)$$

Therefore, $\mathbf{c}^\dagger$ is obtained by determining the signs such that $\zeta_j = \text{sgn}(w_j)$ for components with large absolute values $|w_j|$ of $\mathbf{w} = \mathbf{w}^*(\hat{\mathbf{c}})$, and assigning weights ρ_j within the range where the total amount does not exceed Γ . Thus, the approximate robust regret loss $\tilde{\mathcal{L}}_{\text{rob}}(\hat{\mathbf{c}}, \mathbf{c}_{\text{obs}}) = h_{\mathcal{C}_\Gamma(\mathbf{c}_{\text{obs}})}(\mathbf{w}^*(\hat{\mathbf{c}})) - z(\mathbf{c}^\dagger)$ can be calculated analytically.

Case of Ellipsoidal Uncertainty Set In the case of the ellipsoidal uncertainty set $\mathcal{C}_{\text{ell}}(\mathbf{c}_{\text{obs}})$, the support function is $h_{\mathcal{C}_{\text{ell}}}(\mathbf{w}) = \mathbf{c}_{\text{obs}}^\top \mathbf{w} + \kappa \|\Sigma^{1/2} \mathbf{w}\|_2$, and its maximizing solution $\mathbf{c}^\dagger$ can be analytically calculated as

$$\mathbf{c}^\dagger = \mathbf{c}_{\text{obs}} + \kappa \cdot \frac{\Sigma \mathbf{w}^*(\hat{\mathbf{c}})}{\|\Sigma^{1/2} \mathbf{w}^*(\hat{\mathbf{c}})\|_2} \quad \left(\|\Sigma^{1/2} \mathbf{w}^*(\hat{\mathbf{c}})\|_2 > 0 \right). \quad (25)$$

Therefore, $\tilde{\mathcal{L}}_{\text{rob}}(\hat{\mathbf{c}}, \mathbf{c}_{\text{obs}}) = \mathbf{c}^{\dagger\top} \mathbf{w}^*(\hat{\mathbf{c}}) - z(\mathbf{c}^\dagger)$ can be easily calculated.

B.2 UNIFORM GENERALIZATION OF ROBUST REGRET LOSS

In this section, we show that the robust regret loss defined for observation errors (5) is uniformly bounded from above by the diameter of the feasible set $\mathcal{W}$ and the radius of the uncertainty set $\mathcal{C}(\mathbf{c}_{\text{obs}})$, and derive an upper bound on the generalization error.

B.2.1 Uniform Boundedness of Robust Regret Loss

Preliminaries: Definitions of Diameter and Radius Let $\|\cdot\|$ be a norm and $\|\cdot\|_*$ be its dual norm. We define the diameter of the feasible set $\mathcal{W} \subset \mathbb{R}^n$ with respect to the dual norm as

$$\text{diam}_*(\mathcal{W}) := \sup_{\mathbf{u}, \mathbf{v} \in \mathcal{W}} \|\mathbf{u} - \mathbf{v}\|_* \in [0, \infty). \quad (26)$$

Also, for each $\mathbf{c}_{\text{obs}}$, we define the radius of the uncertainty set $\mathcal{C}(\mathbf{c}_{\text{obs}})$ as

$$R_{\mathcal{C}}(\mathbf{c}_{\text{obs}}) := \sup_{\mathbf{c} \in \mathcal{C}(\mathbf{c}_{\text{obs}})} \|\mathbf{c}\| \in [0, \infty). \quad (27)$$

Hereinafter, we assume the following.

Assumption B.1 (Compactness). $\mathcal{W} \subset \mathbb{R}^n$ is non-empty and compact. For each $\mathbf{c}_{\text{obs}}$, $\mathcal{C}(\mathbf{c}_{\text{obs}}) \subset \mathbb{R}^n$ is non-empty and compact.

From Assumption B.1, $\text{diam}_*(\mathcal{W}) < \infty$ and $R_{\mathcal{C}}(\mathbf{c}_{\text{obs}}) < \infty$ hold.

Upper Bound of Robust Regret Loss

Lemma B.2 (Upper Bound of Robust Regret Loss). Under Assumption B.1, for any $\hat{\mathbf{c}} \in \mathbb{R}^n$ and any $\mathbf{c}_{\text{obs}}$,

$$0 \leq \mathcal{L}_{\text{rob}}(\hat{\mathbf{c}}, \mathbf{c}_{\text{obs}}) \leq R_{\mathcal{C}}(\mathbf{c}_{\text{obs}}) \cdot \text{diam}_*(\mathcal{W}) \quad (28)$$

holds.

(Proof in Section D.6)

Uniform Boundedness in Population In the proof of generalization, it is necessary that the loss is almost surely bounded through the random variable $\tilde{\mathbf{c}}_{\text{obs}}$, so we assume the following.

Assumption B.3 (Uniform Boundedness). *There exists a constant $R < \infty$ such that $R_{\mathcal{C}}(\tilde{\mathbf{c}}_{\text{obs}}) \leq R$ holds almost surely.*

Under Assumption B.3, from Lemma B.2, $0 \leq \mathcal{L}_{\text{rob}}(\hat{\mathbf{c}}, \tilde{\mathbf{c}}_{\text{obs}}) \leq B$ holds almost surely. Here $B := R \cdot \text{diam}_*(\mathcal{W})$.

B.2.2 Generalization Based on Uniform Boundedness

Hereinafter, let $\tilde{\mathbf{s}} := (\tilde{\mathbf{x}}, \tilde{\mathbf{c}}_{\text{obs}}) \sim \mathcal{D}$, and let $S = (\tilde{\mathbf{s}}_1, \dots, \tilde{\mathbf{s}}_N)$ be i.i.d. samples from $\mathcal{D}$. Here $\mathcal{D}$ represents the joint distribution of the covariate $\tilde{\mathbf{x}}$ and the observed coefficient vector $\tilde{\mathbf{c}}_{\text{obs}}$. Let the model class be $\mathcal{F} = \{f_{\boldsymbol{\theta}} : \boldsymbol{\theta} \in \Theta\}$, and the loss class be defined as

$$\mathcal{G} := \left\{ g_{\boldsymbol{\theta}} : (\mathbf{x}, \mathbf{c}_{\text{obs}}) \mapsto \mathcal{L}_{\text{rob}}(f_{\boldsymbol{\theta}}(\mathbf{x}), \mathbf{c}_{\text{obs}}) \mid \boldsymbol{\theta} \in \Theta \right\}. \quad (29)$$

We define the expected loss on the population and the empirical loss as

$$R(\boldsymbol{\theta}) := \mathbb{E}[g_{\boldsymbol{\theta}}(\tilde{\mathbf{s}})], \quad \hat{R}_S(\boldsymbol{\theta}) := \frac{1}{N} \sum_{i=1}^N g_{\boldsymbol{\theta}}(\tilde{\mathbf{s}}_i). \quad (30)$$

Rademacher Complexity We introduce Rademacher complexity as a measure of the expressive power of the loss class $\mathcal{G}$ (how well it can fit random signs on finite samples) [Bartlett and Mendelson, 2002]. First, we define a sequence of independent and identically distributed random variables $\sigma_1, \dots, \sigma_N$ by

$$\Pr(\sigma_i = +1) = \Pr(\sigma_i = -1) = \frac{1}{2}, \quad i = 1, \dots, N. \quad (31)$$

These σ_i are called Rademacher variables.

The empirical Rademacher complexity for sample $S = (\tilde{\mathbf{s}}_1, \dots, \tilde{\mathbf{s}}_N)$ is defined as

$$\hat{\mathfrak{R}}_S(\mathcal{G}) := \mathbb{E}_{\sigma} \left[\sup_{g \in \mathcal{G}} \frac{1}{N} \sum_{i=1}^N \sigma_i g(\tilde{\mathbf{s}}_i) \right]. \quad (32)$$

Here $\mathbb{E}_{\sigma}[\cdot]$ denotes the expectation with respect to the Rademacher variables $\sigma_1, \dots, \sigma_N$.

Furthermore, we define the average Rademacher complexity, which also averages the randomness of the sample itself, as

$$\mathfrak{R}_N(\mathcal{G}) := \mathbb{E}_S[\hat{\mathfrak{R}}_S(\mathcal{G})]. \quad (33)$$

Uniform Generalization

Theorem B.4 (Uniform Convergence for Uniformly Bounded Loss). *Assume Assumption B.3 holds, and thus $0 \leq g_{\boldsymbol{\theta}}(\tilde{\mathbf{s}}) \leq B$ holds almost surely. Then, for any $\delta \in (0, 1)$, with probability at least $1 - \delta$,*

$$\sup_{\boldsymbol{\theta} \in \Theta} (R(\boldsymbol{\theta}) - \hat{R}_S(\boldsymbol{\theta})) \leq 2\hat{\mathfrak{R}}_S(\mathcal{G}) + 3B\sqrt{\frac{\log(2/\delta)}{2N}} \quad (34)$$

holds. Furthermore, taking the expectation,

$$\mathbb{E}_S \left[\sup_{\boldsymbol{\theta} \in \Theta} (R(\boldsymbol{\theta}) - \hat{R}_S(\boldsymbol{\theta})) \right] \leq 2\mathfrak{R}_N(\mathcal{G}) \quad (35)$$

holds.

(Proof in Section D.7)

Corollary B.5 (Generalization Guarantee for Optimal Solution of Robust Regret Loss). *Let $\hat{\boldsymbol{\theta}} \in \arg \min_{\boldsymbol{\theta} \in \Theta} \hat{R}_S(\boldsymbol{\theta})$. From Theorem B.4,*

$$R(\hat{\boldsymbol{\theta}}) \leq \hat{R}_S(\hat{\boldsymbol{\theta}}) + 2\hat{\mathfrak{R}}_S(\mathcal{G}) + 3B\sqrt{\frac{\log(2/\delta)}{2N}}$$

holds.

(Proof in Section D.8)

C DETAILS OF DISTRIBUTIONALLY ROBUST REGRET LOSS

C.1 UNIFORM GENERALIZATION OF DISTRIBUTIONALLY ROBUST REGRET LOSS

C.1.1 Setup

Let $\tilde{s} := (\tilde{x}, \tilde{c}) \sim \mathcal{D}$, and let the sample $S = (\tilde{s}_1, \dots, \tilde{s}_N)$ be i.i.d. samples from $\mathcal{D}$. Here, $\mathcal{D}$ represents the joint distribution of the covariate $\tilde{x}$ and the coefficient vector $\tilde{c}$.

Let the feasible set $\mathcal{W} \subset \mathbb{R}^n$ be non-empty and compact. For the norm $\|\cdot\|$ and the dual norm $\|\cdot\|_*$, we define $\text{diam}_*(\mathcal{W})$ by (26).

We define the optimal value function and the optimal solution as

$$z(c) := \min_{w \in \mathcal{W}} c^\top w, \quad w^*(c) \in \arg \min_{w \in \mathcal{W}} c^\top w.$$

Let a predictive model be $f_\theta : \mathcal{X} \rightarrow \mathbb{R}^n$, and the model class be $\mathcal{F} = \{f_\theta : \theta \in \Theta\}$.

We define the regret as

$$\text{Reg}(\hat{c}, c) := c^\top w^*(\hat{c}) - z(c) \quad (36)$$

.

C.1.2 Definition of Distributionally Robust Regret Loss

Similar to the distributionally robust formulation in the main text, we consider the conditional distribution $P_{\tilde{c}|\tilde{x}=\mathbf{x}}$ of the coefficient vector for each covariate vector $\mathbf{x}$. Here, $\tilde{c}$ represents the random variable of the coefficient vector conditioned on $\tilde{x}$.

For each $\mathbf{x}$, we introduce the 1-Wasserstein ball of the conditional distribution $P_{\tilde{c}|\tilde{x}=\mathbf{x}}$

$$\mathbb{B}_\varepsilon(P_{\tilde{c}|\tilde{x}=\mathbf{x}}) := \left\{ Q_{\tilde{c}|\tilde{x}=\mathbf{x}} \mid W_1(Q_{\tilde{c}|\tilde{x}=\mathbf{x}}, P_{\tilde{c}|\tilde{x}=\mathbf{x}}) \leq \varepsilon \right\}$$

as the uncertainty set. In this case, we define the distributionally robust regret loss against the shift of the conditional distribution as

$$\mathcal{L}_{\text{dro}}(\mathbf{x}; \theta) := \sup_{Q_{\tilde{c}|\tilde{x}=\mathbf{x}} \in \mathbb{B}_\varepsilon(P_{\tilde{c}|\tilde{x}=\mathbf{x}})} \mathbb{E}_{\tilde{c} \sim Q_{\tilde{c}|\tilde{x}=\mathbf{x}}} [\text{Reg}(f_\theta(\mathbf{x}), \tilde{c})] \quad (37)$$

.

Let the population loss be

$$R_{\text{dro}}(\theta) := \mathbb{E}[\mathcal{L}_{\text{dro}}(\tilde{\mathbf{x}}; \theta)] \quad (38)$$

.

By the strong duality argument regarding 1-Wasserstein used in the main text, the right-hand side of (37) can be written in the following form:

$$\mathcal{L}_{\text{dro}}(\mathbf{x}; \theta) = \mathbb{E}[\text{Reg}(f_\theta(\mathbf{x}), \tilde{c}) \mid \tilde{x} = \mathbf{x}] + \varepsilon \text{Lip}(w^*(f_\theta(\mathbf{x}))). \quad (39)$$

Here, $\text{Lip}(w^*(f_\theta(\mathbf{x})))$ is a regularization term depending only on $\mathbf{x}$.

We define the function $g_\theta(\mathbf{x}, c)$ for a sample $(\mathbf{x}, c)$ as follows:

$$g_\theta(\mathbf{x}, c) := \text{Reg}(f_\theta(\mathbf{x}), c) + \varepsilon \text{Lip}(w^*(f_\theta(\mathbf{x}))). \quad (40)$$

Using the law of iterated expectations, the population loss $R_{\text{dro}}(\theta)$ can be expressed as the expectation of g_θ :

$$R_{\text{dro}}(\theta) = \mathbb{E}[\mathcal{L}_{\text{dro}}(\tilde{\mathbf{x}}; \theta)] = \mathbb{E}[\mathbb{E}[g_\theta(\tilde{\mathbf{x}}, \tilde{c}) \mid \tilde{\mathbf{x}}]] = \mathbb{E}[g_\theta(\tilde{\mathbf{x}}, \tilde{c})]. \quad (41)$$

.

Similarly, as an estimator for the expected value of the population loss $R_{\text{dro}}(\boldsymbol{\theta}) = \mathbb{E}[g_{\boldsymbol{\theta}}(\tilde{\mathbf{s}})]$, we define the quantity obtained by replacing the expectation with the empirical distribution (sample mean) as

$$\hat{R}_{\text{dro},S}(\boldsymbol{\theta}) := \frac{1}{N} \sum_{i=1}^N g_{\boldsymbol{\theta}}(\tilde{\mathbf{s}}_i) \quad (42)$$

C.1.3 Uniform Boundedness

Assumption C.1 (Upper Bound on Norm of Coefficient Vector). *There exists a constant $\bar{R}_{\mathcal{C}} < \infty$ such that*

$$\|\tilde{\mathbf{c}}\| \leq \bar{R}_{\mathcal{C}}$$

holds almost surely.

Lemma C.2 (Uniform Boundedness of $g_{\boldsymbol{\theta}}$). *Under Assumption C.1, for any $\boldsymbol{\theta} \in \Theta$ and any $(\mathbf{x}, \mathbf{c})$,*

$$0 \leq g_{\boldsymbol{\theta}}(\mathbf{x}, \mathbf{c}) \leq B_{\text{dro}}$$

holds. Here,

$$B_{\text{dro}} := (\bar{R}_{\mathcal{C}} + \varepsilon) \text{diam}_*(\mathcal{W}). \quad (43)$$

(Proof in Section D.9)

C.1.4 Uniform Generalization of Population Loss

For each $\boldsymbol{\theta} \in \Theta$, the per-sample loss $g_{\boldsymbol{\theta}}: \mathcal{X} \times \mathbb{R}^n \rightarrow \mathbb{R}$ is given by (40). Let the loss class be

$$\mathcal{G}_{\text{dro}} := \{ g_{\boldsymbol{\theta}} \mid \boldsymbol{\theta} \in \Theta \},$$

and define $\hat{\mathfrak{R}}_S(\mathcal{G}_{\text{dro}})$ and $\mathfrak{R}_N(\mathcal{G}_{\text{dro}})$ similarly to (32)–(33).

Theorem C.3 (Uniform Generalization of Population Loss). *Suppose Assumption C.1 holds, and thus $0 \leq g_{\boldsymbol{\theta}}(\tilde{\mathbf{s}}) \leq B_{\text{dro}}$ holds almost surely by Lemma C.2. Then, for any $\delta \in (0, 1)$, with probability at least $1 - \delta$,*

$$\sup_{\boldsymbol{\theta} \in \Theta} (R_{\text{dro}}(\boldsymbol{\theta}) - \hat{R}_{\text{dro},S}(\boldsymbol{\theta})) \leq 2\hat{\mathfrak{R}}_S(\mathcal{G}_{\text{dro}}) + 3B_{\text{dro}} \sqrt{\frac{\log(2/\delta)}{2N}}$$

holds. Furthermore, taking the expectation,

$$\mathbb{E}_S \left[\sup_{\boldsymbol{\theta} \in \Theta} (R_{\text{dro}}(\boldsymbol{\theta}) - \hat{R}_{\text{dro},S}(\boldsymbol{\theta})) \right] \leq 2\mathfrak{R}_N(\mathcal{G}_{\text{dro}})$$

holds.

(Proof in Section D.10)

Corollary C.4 (Generalization Guarantee for Population Loss). *Let $\hat{\boldsymbol{\theta}} \in \arg \min_{\boldsymbol{\theta} \in \Theta} \hat{R}_{\text{dro},S}(\boldsymbol{\theta})$. From Theorem C.3,*

$$R_{\text{dro}}(\hat{\boldsymbol{\theta}}) \leq \hat{R}_{\text{dro},S}(\hat{\boldsymbol{\theta}}) + 2\hat{\mathfrak{R}}_S(\mathcal{G}_{\text{dro}}) + 3B_{\text{dro}} \sqrt{\frac{\log(2/\delta)}{2N}}$$

holds.

(Proof in Section D.11)

C.1.5 Implications of This Section

The properties of the upper bound obtained from the proof in this section can be interpreted as follows.

- The upper bound of g_θ is given by

$$B_{\text{dro}} = (\bar{R}_C + \varepsilon) \text{diam}_*(\mathcal{W})$$

(Lemma C.2). Therefore, the larger the possible magnitude $\bar{R}_C$ of the coefficient vector, the robust radius ε , and the diameter $\text{diam}_*(\mathcal{W})$ of the feasible set, the larger the upper bound B_{dro} becomes, and the larger the second term of the upper bound becomes.

- The gap

$$\sup_{\theta \in \Theta} (R_{\text{dro}}(\theta) - \hat{R}_{\text{dro},S}(\theta))$$

is controlled by $2\hat{\mathfrak{R}}_S(\mathcal{G}_{\text{dro}})$ (Theorem C.3). This is a quantity representing how well the loss class can follow random signs on the sample, meaning that if the expressive power is excessively large, the upper bound becomes loose.

- Increasing ε makes the objective function more conservative in preparation for a wider distribution shift, but since B_{dro} of the upper bound also increases, the gap tends to become large with finite samples.
- Corollary C.4 provides a non-asymptotic guarantee that the learning solution $\hat{\theta}$ that minimizes the DRO objective $\hat{R}_{\text{dro},S}$ based on the empirical distribution also has a small R_{dro} for the unknown population distribution with high probability.

From the above, the upper bound in this section clarifies how the magnitude of uncertainty (radius), the geometry of optimization (diameter), and the model expressive power affect generalization.

C.2 GRADIENT CALCULATION OF DISTRIBUTIONALLY ROBUST LOSS

We describe the details of the gradient calculation of the distributionally robust loss. Regarding $\text{Lip}(\mathbf{w}^*(\hat{c}))$, we need to calculate the gradient through the subgradient of

$$\text{Lip}(\mathbf{w}^*(\hat{c})) = \sup_{\mathbf{w} \in \mathcal{W}} \|\mathbf{w}^*(\hat{c}) - \mathbf{w}\|_*$$

introduced in the previous section.

To do this, using the dual norm

$$\|\mathbf{u}\|_* = \max_{\mathbf{g} \in \mathcal{G}} \mathbf{g}^\top \mathbf{u}, \quad \mathcal{G} := \{\mathbf{g} \in \mathbb{R}^n : \|\mathbf{g}\| \leq 1\},$$

we rewrite it as

$$\text{Lip}(\mathbf{w}^*(\hat{c})) = \max_{(\mathbf{w}, \mathbf{g}) \in \mathcal{W} \times \mathcal{G}} \mathbf{g}^\top (\mathbf{w}^*(\hat{c}) - \mathbf{w}).$$

Here, $\mathcal{V} := \mathcal{W} \times \mathcal{G}$ is a compact set, and letting $\phi(\mathbf{w}^*(\hat{c}); \mathbf{w}, \mathbf{g}) := \mathbf{g}^\top (\mathbf{w}^*(\hat{c}) - \mathbf{w})$, ϕ is linear and continuous with respect to $(\mathbf{w}^*(\hat{c}), \mathbf{w}, \mathbf{g})$, and for each $(\mathbf{w}, \mathbf{g})$, $\mathbf{w}^*(\hat{c}) \mapsto \phi(\mathbf{w}^*(\hat{c}); \mathbf{w}, \mathbf{g})$ is C^1 and convex, and

$$\nabla_{\mathbf{w}^*(\hat{c})} \phi(\mathbf{w}^*(\hat{c}); \mathbf{w}, \mathbf{g}) = \mathbf{g}$$

holds. Furthermore, $(\mathbf{w}, \mathbf{g}) \mapsto \nabla_{\mathbf{w}^*(\hat{c})} \phi(\mathbf{w}^*(\hat{c}); \mathbf{w}, \mathbf{g}) = \mathbf{g}$ is continuous on $\mathcal{V}$. Therefore, the assumptions of Danskin's theorem are satisfied.

Thus, letting

$$\text{Lip}(\mathbf{w}^*(\hat{c})) = \max_{v \in \mathcal{V}} \phi(\mathbf{w}^*(\hat{c}); v),$$

we have

$$\partial \text{Lip}(\mathbf{w}^*(\hat{c})) = \text{conv}\{\nabla_{\mathbf{w}^*(\hat{c})} \phi(\mathbf{w}^*(\hat{c}); v) \mid v \in V(\mathbf{w}^*(\hat{c}))\}.$$

The gradient of the parameters θ of the predictive model with respect to the distributionally robust loss $\mathcal{L}_{\text{dro}}(\mathbf{x})$ can be written by the chain rule as

$$\frac{\partial}{\partial \hat{c}} \text{Lip}(\mathbf{w}^*(\hat{c})) = J_{\mathbf{w}^*(\hat{c})}^\top \mathbf{g}, \quad \mathbf{g} \in \partial \text{Lip}(\mathbf{w}^*(\hat{c})).$$

Here, $J_{\mathbf{w}^*(\hat{\mathbf{c}})}$ is the Jacobian matrix of $\mathbf{w}^*(\hat{\mathbf{c}})$, and $\mathbf{g}$ is the subgradient obtained from Danskin's theorem above. In implementation, we perform the calculation using methods such as the one proposed in [Huang and Gupta, 2024] as a method to calculate the approximate gradient of $\mathbf{w}^*(\hat{\mathbf{c}})$ without using the gradient of SPO+.

D PROOFS

D.1 PROOF OF THEOREM 3.1

First, as a preparation for gradient calculation, we show a lemma about the continuity of the optimal value function $z(\mathbf{c})$.

Lemma D.1. *Let $\mathcal{W} \subseteq \mathbb{R}^n$ be a compact set, and define $z : \mathbb{R}^n \rightarrow \mathbb{R}$ by $z(\mathbf{c}) := \min_{\mathbf{w} \in \mathcal{W}} \mathbf{c}^\top \mathbf{w}$. Then z is Lipschitz continuous.*

The proof of this lemma is given in the next section (Section D.2).

Proof of Theorem 3.1. We apply Danskin's theorem [Bertsekas, 1999, Proposition B.25] (Proposition A.1). For the proposed robust regret loss, regarding $\mathbf{w}^*(\hat{\mathbf{c}})$ as a variable and $\mathbf{c} \in \mathcal{C}$ as a coefficient vector, let

$$\psi(\mathbf{w}^*(\hat{\mathbf{c}}), \mathbf{c}) := \mathbf{c}^\top \mathbf{w}^*(\hat{\mathbf{c}}) - z(\mathbf{c}).$$

From the non-empty compactness of $\mathcal{W}, \mathcal{C}$ and the continuity of $z(\mathbf{c})$ (Lemma D.1), ψ is continuous in $(\mathbf{w}^*(\hat{\mathbf{c}}), \mathbf{c})$, and since it is affine with respect to $\mathbf{w}^*(\hat{\mathbf{c}})$ for each $\mathbf{c}$, it is convex and C^1 , and

$$\nabla_{\mathbf{w}^*(\hat{\mathbf{c}})} \psi(\mathbf{w}^*(\hat{\mathbf{c}}), \mathbf{c}) = \mathbf{c}.$$

Also, since $\mathcal{C}$ is compact, it satisfies the assumption of Proposition A.1. Therefore, the subgradient set of the robust regret loss with respect to $\mathbf{w}^*(\hat{\mathbf{c}})$ is derived as the convex hull of the worst-case coefficient vectors $\mathbf{c}^* \in \arg \max_{\mathbf{c} \in \mathcal{C}} \psi(\mathbf{w}^*(\hat{\mathbf{c}}), \mathbf{c})$:

$$\partial_{\mathbf{w}^*(\hat{\mathbf{c}})} \mathcal{L}_{\text{rob}}(\hat{\mathbf{c}}, \mathbf{c}_{\text{obs}}) = \text{conv} \left\{ \arg \max_{\mathbf{c} \in \mathcal{C}(\mathbf{c}_{\text{obs}})} (\mathbf{c}^\top \mathbf{w}^*(\hat{\mathbf{c}}) - z(\mathbf{c})) \right\}. \quad (44)$$

□

D.2 PROOF OF LEMMA D.1

Proof. From the compactness of $\mathcal{W}$,

$$M := \sup_{\mathbf{w} \in \mathcal{W}} \|\mathbf{w}\| < \infty$$

holds. Take arbitrary $\mathbf{c}, \mathbf{c}' \in \mathbb{R}^n$, and let $\mathbf{w}' \in \mathcal{W}$ be a point satisfying $z(\mathbf{c}') = \min_{\mathbf{w} \in \mathcal{W}} \mathbf{c}'^\top \mathbf{w} = \mathbf{c}'^\top \mathbf{w}'$. Note that since $\mathcal{W}$ is compact and the objective function is continuous, the minimum is always achieved at some point in $\mathcal{W}$. At this time,

$$\begin{aligned} z(\mathbf{c}) - z(\mathbf{c}') &= z(\mathbf{c}) - \mathbf{c}'^\top \mathbf{w}' \\ &\leq \mathbf{c}^\top \mathbf{w}' - \mathbf{c}'^\top \mathbf{w}' \\ &= (\mathbf{c} - \mathbf{c}')^\top \mathbf{w}' \\ &\leq \|\mathbf{c} - \mathbf{c}'\| \|\mathbf{w}'\| \\ &\leq M \|\mathbf{c} - \mathbf{c}'\| \end{aligned}$$

holds. Similarly, by swapping $\mathbf{c}$ and $\mathbf{c}'$,

$$z(\mathbf{c}') - z(\mathbf{c}) \leq M \|\mathbf{c} - \mathbf{c}'\|$$

also follows, so

$$|z(\mathbf{c}) - z(\mathbf{c}')| \leq M \|\mathbf{c} - \mathbf{c}'\|$$

holds for arbitrary $\mathbf{c}, \mathbf{c}' \in \mathbb{R}^n$. Therefore, z is Lipschitz continuous. □

D.3 PROOF OF THEOREM 3.2

Proof. To derive an equivalent representation of the distributionally robust loss $\mathcal{L}_{\text{dro}}(\mathbf{x})$, we apply the Blanchet–Murthy strong duality theorem [Blanchet and Murthy, 2019]. As a preliminary step, we first state a proposition about the Lipschitz continuity of the objective function $\phi_{\mathbf{w}^*(\hat{\mathbf{c}})}$.

Proposition D.2 (Lipschitz continuity of $\phi_{\mathbf{w}^*(\hat{\mathbf{c}})}$). *Let the constant for arbitrary $\mathbf{c}, \mathbf{c}' \in \mathbb{R}^n$ be*

$$\text{Lip}(\mathbf{w}^*(\hat{\mathbf{c}})) := \sup_{\mathbf{w} \in \mathcal{W}} \|\mathbf{w}^*(\hat{\mathbf{c}}) - \mathbf{w}\|_*,$$

then the following inequality holds:

$$|\phi_{\mathbf{w}^*(\hat{\mathbf{c}})}(\mathbf{c}) - \phi_{\mathbf{w}^*(\hat{\mathbf{c}})}(\mathbf{c}')| \leq \text{Lip}(\mathbf{w}^*(\hat{\mathbf{c}})) \|\mathbf{c} - \mathbf{c}'\|.$$

That is, $\phi_{\mathbf{w}^(\hat{\mathbf{c}})}$ is Lipschitz continuous with respect to the norm $\|\cdot\|$, and $\text{Lip}(\mathbf{w}^*(\hat{\mathbf{c}}))$ is the Lipschitz constant.*

The proof of this proposition is given in the next section (Section D.4).

From Proposition D.2, $\phi_{\mathbf{w}^*(\hat{\mathbf{c}})}$ is a Lipschitz continuous function and is integrable with respect to $P_{\tilde{\mathbf{c}}|\mathbf{x}}$.

Letting the distance be $d(\mathbf{u}, \mathbf{v}) = \|\mathbf{u} - \mathbf{v}\|$, Assumptions (A1), (A2) of the Blanchet–Murthy strong duality [Blanchet and Murthy, 2019] are satisfied by Proposition D.2. Therefore,

$$\sup_{Q_{\tilde{\mathbf{c}}|\mathbf{x}} \in \mathbb{B}_\varepsilon(P_{\tilde{\mathbf{c}}|\mathbf{x}})} \mathbb{E}_{Q_{\tilde{\mathbf{c}}|\mathbf{x}}}[\phi_{\mathbf{w}^*(\hat{\mathbf{c}})}(\tilde{\mathbf{c}})] = \inf_{\lambda \geq 0} \left\{ \lambda \varepsilon + \mathbb{E}_{P_{\tilde{\mathbf{c}}|\mathbf{x}}} \left[\sup_{\mathbf{y}} \{ \phi_{\mathbf{w}^*(\hat{\mathbf{c}})}(\mathbf{y}) - \lambda \|\mathbf{y} - \tilde{\mathbf{c}}\| \} \right] \right\}$$

holds. Although maximization by $\mathbf{y}$ for each $\tilde{\mathbf{c}}$ remains on the right side, by the following Theorem D.3, this term is divided into a finite value or infinity depending on the value of λ .

Theorem D.3. *For any $\lambda \geq 0$, $\tilde{\mathbf{c}} \in \mathbb{R}^n$,*

$$\sup_{\mathbf{y} \in \mathbb{R}^n} \{ \phi_{\mathbf{w}^*(\hat{\mathbf{c}})}(\mathbf{y}) - \lambda \|\mathbf{y} - \tilde{\mathbf{c}}\| \} = \begin{cases} \phi_{\mathbf{w}^*(\hat{\mathbf{c}})}(\tilde{\mathbf{c}}), & \lambda \geq \text{Lip}(\mathbf{w}^*(\hat{\mathbf{c}})), \\ +\infty, & \lambda < \text{Lip}(\mathbf{w}^*(\hat{\mathbf{c}})), \end{cases}$$

holds. In particular, the optimal value in

$$\inf_{\lambda \geq 0} \left\{ \lambda \varepsilon + \mathbb{E}_{\mu} \left[\sup_{\mathbf{y}} \{ \phi_{\mathbf{w}^*(\hat{\mathbf{c}})}(\mathbf{y}) - \lambda \|\mathbf{y} - \tilde{\mathbf{c}}\| \} \right] \right\}$$

appearing in the Blanchet–Murthy strong duality is achieved at $\lambda^ = \text{Lip}(\mathbf{w}^*(\hat{\mathbf{c}}))$.*

The proof of this theorem is given in Section D.5.

Using Theorem D.3, the optimal value of the dual problem is achieved at $\lambda = \text{Lip}(\mathbf{w}^*(\hat{\mathbf{c}}))$, and

$$\mathcal{L}_{\text{dro}}(\mathbf{x}) = \mathbb{E}_{\tilde{\mathbf{c}} \sim P_{\tilde{\mathbf{c}}|\mathbf{x}}} [\phi_{\mathbf{w}^*(\hat{\mathbf{c}})}(\tilde{\mathbf{c}})] + \varepsilon \text{Lip}(\mathbf{w}^*(\hat{\mathbf{c}}))$$

is obtained. □

D.4 PROOF OF PROPOSITION D.2

Proof. $\mathcal{W}$ is compact, and since the mapping $\mathbf{w} \mapsto \mathbf{c}^\top (\mathbf{w}^*(\hat{\mathbf{c}}) - \mathbf{w})$ is continuous, the maximum value is always achieved at some point in $\mathcal{W}$. For each $\mathbf{c} \in \mathbb{R}^n$, choose one

$$\mathbf{w}_{\mathbf{c}} \in \arg \max_{\mathbf{w} \in \mathcal{W}} \mathbf{c}^\top (\mathbf{w}^*(\hat{\mathbf{c}}) - \mathbf{w}).$$

Then

$$\phi_{\mathbf{w}^*(\hat{\mathbf{c}})}(\mathbf{c}) = \mathbf{c}^\top (\mathbf{w}^*(\hat{\mathbf{c}}) - \mathbf{w}_{\mathbf{c}}) \leq \mathbf{c}^\top (\mathbf{w}^*(\hat{\mathbf{c}}) - \mathbf{w}_{\mathbf{c}'})$$

holds, so

$$\begin{aligned}\phi_{\mathbf{w}^*(\hat{\mathbf{c}})}(\mathbf{c}) - \phi_{\mathbf{w}^*(\hat{\mathbf{c}})}(\mathbf{c}') &\leq \mathbf{c}^\top (\mathbf{w}^*(\hat{\mathbf{c}}) - \mathbf{w}_{\mathbf{c}'}) - \mathbf{c}'^\top (\mathbf{w}^*(\hat{\mathbf{c}}) - \mathbf{w}_{\mathbf{c}'}) \\ &= (\mathbf{c} - \mathbf{c}')^\top (\mathbf{w}^*(\hat{\mathbf{c}}) - \mathbf{w}_{\mathbf{c}'}).\end{aligned}$$

From the definition of the dual norm

$$(\mathbf{c} - \mathbf{c}')^\top (\mathbf{w}^*(\hat{\mathbf{c}}) - \mathbf{w}_{\mathbf{c}'}) \leq \|\mathbf{c} - \mathbf{c}'\| \|\mathbf{w}^*(\hat{\mathbf{c}}) - \mathbf{w}_{\mathbf{c}'}\|_* \leq \|\mathbf{c} - \mathbf{c}'\| \sup_{\mathbf{w} \in \mathcal{W}} \|\mathbf{w}^*(\hat{\mathbf{c}}) - \mathbf{w}\|_* = \text{Lip}(\mathbf{w}^*(\hat{\mathbf{c}})) \|\mathbf{c} - \mathbf{c}'\|.$$

Therefore

$$\phi_{\mathbf{w}^*(\hat{\mathbf{c}})}(\mathbf{c}) - \phi_{\mathbf{w}^*(\hat{\mathbf{c}})}(\mathbf{c}') \leq \text{Lip}(\mathbf{w}^*(\hat{\mathbf{c}})) \|\mathbf{c} - \mathbf{c}'\|.$$

Similarly, swapping the roles of $\mathbf{c}$ and $\mathbf{c}'$,

$$\phi_{\mathbf{w}^*(\hat{\mathbf{c}})}(\mathbf{c}') - \phi_{\mathbf{w}^*(\hat{\mathbf{c}})}(\mathbf{c}) \leq \text{Lip}(\mathbf{w}^*(\hat{\mathbf{c}})) \|\mathbf{c} - \mathbf{c}'\|$$

also follows, so combining the two,

$$|\phi_{\mathbf{w}^*(\hat{\mathbf{c}})}(\mathbf{c}) - \phi_{\mathbf{w}^*(\hat{\mathbf{c}})}(\mathbf{c}')| \leq \text{Lip}(\mathbf{w}^*(\hat{\mathbf{c}})) \|\mathbf{c} - \mathbf{c}'\|$$

holds for arbitrary $\mathbf{c}, \mathbf{c}' \in \mathbb{R}^n$. This shows the claim. $\square$

D.5 PROOF OF THEOREM D.3

Proof. First, from Proposition D.2, $\phi_{\mathbf{w}^*(\hat{\mathbf{c}})}$ satisfies

$$|\phi_{\mathbf{w}^*(\hat{\mathbf{c}})}(\mathbf{y}) - \phi_{\mathbf{w}^*(\hat{\mathbf{c}})}(\tilde{\mathbf{c}})| \leq \text{Lip}(\mathbf{w}^*(\hat{\mathbf{c}})) \|\mathbf{y} - \tilde{\mathbf{c}}\| \quad (\forall \mathbf{y}, \tilde{\mathbf{c}} \in \mathbb{R}^n)$$

with respect to the norm $\|\cdot\|$.

(1) Case $\lambda \geq \text{Lip}(\mathbf{w}^*(\hat{\mathbf{c}}))$: For any $\mathbf{y} \in \mathbb{R}^n$,

$$\phi_{\mathbf{w}^*(\hat{\mathbf{c}})}(\mathbf{y}) \leq \phi_{\mathbf{w}^*(\hat{\mathbf{c}})}(\tilde{\mathbf{c}}) + \text{Lip}(\mathbf{w}^*(\hat{\mathbf{c}})) \|\mathbf{y} - \tilde{\mathbf{c}}\| \leq \phi_{\mathbf{w}^*(\hat{\mathbf{c}})}(\tilde{\mathbf{c}}) + \lambda \|\mathbf{y} - \tilde{\mathbf{c}}\|$$

holds. Therefore

$$\phi_{\mathbf{w}^*(\hat{\mathbf{c}})}(\mathbf{y}) - \lambda \|\mathbf{y} - \tilde{\mathbf{c}}\| \leq \phi_{\mathbf{w}^*(\hat{\mathbf{c}})}(\tilde{\mathbf{c}}) \quad (\forall \mathbf{y}).$$

On the other hand, if we take $\mathbf{y} = \tilde{\mathbf{c}}$, the left side becomes $\phi_{\mathbf{w}^*(\hat{\mathbf{c}})}(\tilde{\mathbf{c}}) - \lambda \|\tilde{\mathbf{c}} - \tilde{\mathbf{c}}\| = \phi_{\mathbf{w}^*(\hat{\mathbf{c}})}(\tilde{\mathbf{c}})$, so the following can be derived:

$$\sup_{\mathbf{y} \in \mathbb{R}^n} \{\phi_{\mathbf{w}^*(\hat{\mathbf{c}})}(\mathbf{y}) - \lambda \|\mathbf{y} - \tilde{\mathbf{c}}\|\} = \phi_{\mathbf{w}^*(\hat{\mathbf{c}})}(\tilde{\mathbf{c}}).$$

(2) Case $\lambda < \text{Lip}(\mathbf{w}^*(\hat{\mathbf{c}}))$: From the definition of $L(\mathbf{w}^*(\hat{\mathbf{c}}))$,

$$\text{Lip}(\mathbf{w}^*(\hat{\mathbf{c}})) = \sup_{\mathbf{w} \in \mathcal{W}} \|\mathbf{w}^*(\hat{\mathbf{c}}) - \mathbf{w}\|_*.$$

Since $\mathcal{W}$ is a compact set and the mapping

$$\mathbf{w} \mapsto \|\mathbf{w}^*(\hat{\mathbf{c}}) - \mathbf{w}\|_*$$

is continuous, there exists a point $\mathbf{w}^* \in \mathcal{W}$ achieving the maximum value, and

$$\|\mathbf{w}^*(\hat{\mathbf{c}}) - \mathbf{w}^*\|_* = \text{Lip}(\mathbf{w}^*(\hat{\mathbf{c}}))$$

holds.

Next, if we take the unit sphere

$$\mathcal{G} := \{\mathbf{g} \in \mathbb{R}^n : \|\mathbf{g}\| = 1\},$$

since $\mathcal{G}$ is compact and the mapping

$$\mathbf{g} \mapsto \langle \mathbf{g}, \mathbf{w}^*(\hat{\mathbf{c}}) - \mathbf{w}^* \rangle$$

is continuous, there exists $\mathbf{g}^* \in \mathcal{G}$ achieving the maximum value, satisfying

$$\max_{\mathbf{g} \in \mathcal{G}} \mathbf{g}^\top (\mathbf{w}^*(\hat{\mathbf{c}}) - \mathbf{w}^*) = \|\mathbf{w}^*(\hat{\mathbf{c}}) - \mathbf{w}^*\|_* = \text{Lip}(\mathbf{w}^*(\hat{\mathbf{c}})).$$

For arbitrary $t > 0$, let

$$\mathbf{y}_t := \tilde{\mathbf{c}} + t \mathbf{g}^*.$$

From the definition of $\phi_{\mathbf{w}^*(\hat{\mathbf{c}})}$

$$\phi_{\mathbf{w}^*(\hat{\mathbf{c}})}(\mathbf{y}) = \max_{\mathbf{w} \in \mathcal{W}} \mathbf{y}^\top (\mathbf{w}^*(\hat{\mathbf{c}}) - \mathbf{w}),$$

we have

$$\phi_{\mathbf{w}^*(\hat{\mathbf{c}})}(\mathbf{y}_t) \geq \mathbf{y}_t^\top (\mathbf{w}^*(\hat{\mathbf{c}}) - \mathbf{w}^*) = \tilde{\mathbf{c}}^\top (\mathbf{w}^*(\hat{\mathbf{c}}) - \mathbf{w}^*) + t (\mathbf{g}^{*\top} (\mathbf{w}^*(\hat{\mathbf{c}}) - \mathbf{w}^*)) = \tilde{\mathbf{c}}^\top (\mathbf{w}^*(\hat{\mathbf{c}}) - \mathbf{w}^*) + t \text{Lip}(\mathbf{w}^*(\hat{\mathbf{c}})).$$

On the other hand, since $\|\mathbf{y}_t - \tilde{\mathbf{c}}\| = \|t \mathbf{v}^*\| = t$,

$$\phi_{\mathbf{w}^*(\hat{\mathbf{c}})}(\mathbf{y}_t) - \lambda \|\mathbf{y}_t - \tilde{\mathbf{c}}\| \geq \tilde{\mathbf{c}}^\top (\mathbf{w}^*(\hat{\mathbf{c}}) - \mathbf{w}^*) + t(\text{Lip}(\mathbf{w}^*(\hat{\mathbf{c}})) - \lambda).$$

Here, since $\lambda < \text{Lip}(\mathbf{w}^*(\hat{\mathbf{c}}))$, $\text{Lip}(\mathbf{w}^*(\hat{\mathbf{c}})) - \lambda > 0$, and the right side diverges to $+\infty$ as $t \rightarrow \infty$. Therefore

$$\sup_{\mathbf{y} \in \mathbb{R}^n} \{\phi_{\mathbf{w}^*(\hat{\mathbf{c}})}(\mathbf{y}) - \lambda \|\mathbf{y} - \tilde{\mathbf{c}}\|\} = +\infty$$

follows.

From the above, for arbitrary $\lambda \geq 0$, $\tilde{\mathbf{c}} \in \mathbb{R}^n$, it is shown that

$$\sup_{\mathbf{y} \in \mathbb{R}^n} \{\phi_{\mathbf{w}^*(\hat{\mathbf{c}})}(\mathbf{y}) - \lambda \|\mathbf{y} - \tilde{\mathbf{c}}\|\} = \begin{cases} \phi_{\mathbf{w}^*(\hat{\mathbf{c}})}(\tilde{\mathbf{c}}), & \lambda \geq \text{Lip}(\mathbf{w}^*(\hat{\mathbf{c}})), \\ +\infty, & \lambda < \text{Lip}(\mathbf{w}^*(\hat{\mathbf{c}})) \end{cases}$$

holds. Finally, in the dual problem

$$\inf_{\lambda \geq 0} \left\{ \lambda \varepsilon + \mathbb{E}_\mu \left[\sup_{\mathbf{y}} \{\phi_{\mathbf{w}^*(\hat{\mathbf{c}})}(\mathbf{y}) - \lambda \|\mathbf{y} - \tilde{\mathbf{c}}\|\} \right] \right\},$$

in the range $\lambda < \text{Lip}(\mathbf{w}^*(\hat{\mathbf{c}}))$, the right side becomes $+\infty$, and for $\lambda \geq \text{Lip}(\mathbf{w}^*(\hat{\mathbf{c}}))$, the right side is $\lambda \varepsilon + \mathbb{E}_\mu[\phi_{\mathbf{w}^*(\hat{\mathbf{c}})}(\tilde{\mathbf{c}})]$, so the minimum value with respect to λ is achieved at $\lambda^* = \text{Lip}(\mathbf{w}^*(\hat{\mathbf{c}}))$. $\square$

D.6 PROOF OF LEMMA B.2

Proof. First, we show non-negativity. Fix arbitrary $\mathbf{c} \in \mathcal{C}(\mathbf{c}_{\text{obs}})$. From $z(\mathbf{c}) = \min_{\mathbf{w} \in \mathcal{W}} \mathbf{c}^\top \mathbf{w}$, for any $\mathbf{w} \in \mathcal{W}$,

$$z(\mathbf{c}) \leq \mathbf{c}^\top \mathbf{w} \tag{45}$$

holds. In particular, taking $\mathbf{w} = \mathbf{w}^*(\hat{\mathbf{c}}) \in \mathcal{W}$,

$$\mathbf{c}^\top \mathbf{w}^*(\hat{\mathbf{c}}) - z(\mathbf{c}) \geq \mathbf{c}^\top \mathbf{w}^*(\hat{\mathbf{c}}) - \mathbf{c}^\top \mathbf{w}^*(\hat{\mathbf{c}}) = 0.$$

Therefore, even taking the maximum,

$$\mathcal{L}_{\text{rob}}(\hat{\mathbf{c}}, \mathbf{c}_{\text{obs}}) = \max_{\mathbf{c} \in \mathcal{C}(\mathbf{c}_{\text{obs}})} \{\mathbf{c}^\top \mathbf{w}^*(\hat{\mathbf{c}}) - z(\mathbf{c})\} \geq 0$$

follows.

Next, we show the upper bound. For arbitrary $\mathbf{c} \in \mathcal{C}(\mathbf{c}_{\text{obs}})$,

$$\begin{aligned} \mathbf{c}^\top \mathbf{w}^*(\hat{\mathbf{c}}) - z(\mathbf{c}) &= \mathbf{c}^\top \mathbf{w}^*(\hat{\mathbf{c}}) - \min_{\mathbf{w} \in \mathcal{W}} \mathbf{c}^\top \mathbf{w} \\ &= \max_{\mathbf{w} \in \mathcal{W}} (\mathbf{c}^\top \mathbf{w}^*(\hat{\mathbf{c}}) - \mathbf{c}^\top \mathbf{w}) \\ &= \max_{\mathbf{w} \in \mathcal{W}} \mathbf{c}^\top (\mathbf{w}^*(\hat{\mathbf{c}}) - \mathbf{w}). \end{aligned}$$

From Hölder's inequality, for arbitrary $\mathbf{w} \in \mathcal{W}$,

$$\mathbf{c}^\top (\mathbf{w}^*(\hat{\mathbf{c}}) - \mathbf{w}) \leq \|\mathbf{c}\| \cdot \|\mathbf{w}^*(\hat{\mathbf{c}}) - \mathbf{w}\|_* \quad (46)$$

holds. Therefore

$$\max_{\mathbf{w} \in \mathcal{W}} \mathbf{c}^\top (\mathbf{w}^*(\hat{\mathbf{c}}) - \mathbf{w}) \leq \|\mathbf{c}\| \cdot \max_{\mathbf{w} \in \mathcal{W}} \|\mathbf{w}^*(\hat{\mathbf{c}}) - \mathbf{w}\|_*.$$

Here, since $\mathbf{w}^*(\hat{\mathbf{c}}) \in \mathcal{W}$,

$$\max_{\mathbf{w} \in \mathcal{W}} \|\mathbf{w}^*(\hat{\mathbf{c}}) - \mathbf{w}\|_* \leq \sup_{\mathbf{u}, \mathbf{v} \in \mathcal{W}} \|\mathbf{u} - \mathbf{v}\|_* = \text{diam}_*(\mathcal{W}).$$

Therefore, for arbitrary $\mathbf{c} \in \mathcal{C}(\mathbf{c}_{\text{obs}})$,

$$\mathbf{c}^\top \mathbf{w}^*(\hat{\mathbf{c}}) - z(\mathbf{c}) \leq \|\mathbf{c}\| \cdot \text{diam}_*(\mathcal{W})$$

is obtained. Taking the maximum of $\mathbf{c} \in \mathcal{C}(\mathbf{c}_{\text{obs}})$ on both sides,

$$\mathcal{L}_{\text{rob}}(\hat{\mathbf{c}}, \mathbf{c}_{\text{obs}}) \leq \left(\sup_{\mathbf{c} \in \mathcal{C}(\mathbf{c}_{\text{obs}})} \|\mathbf{c}\| \right) \cdot \text{diam}_*(\mathcal{W}) = R_{\mathcal{C}}(\mathbf{c}_{\text{obs}}) \cdot \text{diam}_*(\mathcal{W})$$

becomes. Thus the claim holds. $\square$

D.7 PROOF OF THEOREM B.4

Proof. We show the expected value bound (35) and generalization (34) in order.

Expected Value Bound Let $\mathcal{D} = (\tilde{s}_1, \dots, \tilde{s}_N)$ and $\mathcal{D}' = (\tilde{s}'_1, \dots, \tilde{s}'_N)$ be i.i.d. draws from $\mathcal{P}$. For each θ ,

$$R(\theta) = \mathbb{E}[g_\theta(\tilde{s})] = \mathbb{E}_{\mathcal{D}'} \left[\frac{1}{N} \sum_{i=1}^N g_\theta(\tilde{s}'_i) \right] = \mathbb{E}_{\mathcal{D}'} [\hat{R}_{\mathcal{D}'}(\theta)].$$

Therefore

$$R(\theta) - \hat{R}_{\mathcal{D}}(\theta) = \mathbb{E}_{\mathcal{D}'} [\hat{R}_{\mathcal{D}'}(\theta)] - \hat{R}_{\mathcal{D}}(\theta).$$

Taking $\sup_\theta$ and expectation with respect to $\mathcal{D}$, by the convexity of $\sup$ and Jensen's inequality,

$$\begin{aligned} \mathbb{E}_{\mathcal{D}} \left[\sup_{\theta} (R(\theta) - \hat{R}_{\mathcal{D}}(\theta)) \right] &= \mathbb{E}_{\mathcal{D}} \left[\sup_{\theta} (\mathbb{E}_{\mathcal{D}'} [\hat{R}_{\mathcal{D}'}(\theta)] - \hat{R}_{\mathcal{D}}(\theta)) \right] \\ &\leq \mathbb{E}_{\mathcal{D}, \mathcal{D}'} \left[\sup_{\theta} (\hat{R}_{\mathcal{D}'}(\theta) - \hat{R}_{\mathcal{D}}(\theta)) \right] \\ &= \mathbb{E}_{\mathcal{D}, \mathcal{D}'} \left[\sup_{\theta} \frac{1}{N} \sum_{i=1}^N (g_\theta(\tilde{s}'_i) - g_\theta(\tilde{s}_i)) \right]. \end{aligned}$$

Here, we introduce independent Rademacher variables $\sigma_1, \dots, \sigma_N$. Since $\mathcal{D}$ and $\mathcal{D}'$ are identically distributed and independent, the joint distribution of $\{(\tilde{s}_i, \tilde{s}'_i)\}_{i=1}^N$ does not change even if $\tilde{s}_i$ and $\tilde{s}'_i$ are swapped for each i . Due to this symmetry, the expected value does not change even if the sign of the difference $(g_\theta(\tilde{s}'_i) - g_\theta(\tilde{s}_i))$ is randomly flipped by $\sigma_i \in \{+1, -1\}$, and

$$\mathbb{E}_{\mathcal{D}, \mathcal{D}'} \left[\sup_{\theta} \frac{1}{N} \sum_{i=1}^N (g_\theta(\tilde{s}'_i) - g_\theta(\tilde{s}_i)) \right] = \mathbb{E}_{\mathcal{D}, \mathcal{D}', \sigma} \left[\sup_{\theta} \frac{1}{N} \sum_{i=1}^N \sigma_i (g_\theta(\tilde{s}'_i) - g_\theta(\tilde{s}_i)) \right]$$

holds.

Also, since $\sup_{\theta} (a_{\theta} + b_{\theta}) \leq \sup_{\theta} a_{\theta} + \sup_{\theta} b_{\theta}$,

$$\sup_{\theta} \sum_{i=1}^N \sigma_i (g_\theta(\tilde{s}'_i) - g_\theta(\tilde{s}_i)) \leq \sup_{\theta} \sum_{i=1}^N \sigma_i g_\theta(\tilde{s}'_i) + \sup_{\theta} \sum_{i=1}^N (-\sigma_i) g_\theta(\tilde{s}_i).$$

Therefore

$$\mathbb{E}_{\mathcal{D}} \left[\sup_{\boldsymbol{\theta}} (R(\boldsymbol{\theta}) - \widehat{R}_{\mathcal{D}}(\boldsymbol{\theta})) \right] \leq \mathbb{E}_{\mathcal{D}', \sigma} \left[\sup_{\boldsymbol{\theta}} \frac{1}{N} \sum_i \sigma_i g_{\boldsymbol{\theta}}(\tilde{s}'_i) \right] + \mathbb{E}_{\mathcal{D}, \sigma} \left[\sup_{\boldsymbol{\theta}} \frac{1}{N} \sum_i (-\sigma_i) g_{\boldsymbol{\theta}}(\tilde{s}_i) \right].$$

Since Rademacher variables are identically distributed even if signs are flipped, the two terms are equal, and

$$\mathbb{E}_{\mathcal{D}} \left[\sup_{\boldsymbol{\theta}} (R(\boldsymbol{\theta}) - \widehat{R}_{\mathcal{D}}(\boldsymbol{\theta})) \right] \leq 2 \mathbb{E}_{\mathcal{D}, \sigma} \left[\sup_{g \in \mathcal{G}} \frac{1}{N} \sum_i \sigma_i g(\tilde{s}_i) \right] = 2 \mathfrak{R}_N(\mathcal{G})$$

is obtained. This shows (35).

Generalization First, we directly evaluate the amount of change in $\Phi(\mathcal{D})$ due to replacing one sample point. Let

$$\Phi(\mathcal{D}) := \sup_{\boldsymbol{\theta} \in \Theta} (R(\boldsymbol{\theta}) - \widehat{R}_{\mathcal{D}}(\boldsymbol{\theta})).$$

Here $\mathcal{D} = (\tilde{s}_1, \dots, \tilde{s}_N)$ are i.i.d. samples. Let $\mathcal{D}^{(i)}$ be the sample obtained by replacing the i -th element $\tilde{s}_i$ of $\mathcal{D}$ with an independent identically distributed sample $\tilde{s}'_i$:

$$\mathcal{D}^{(i)} := (\tilde{s}_1, \dots, \tilde{s}_{i-1}, \tilde{s}'_i, \tilde{s}_{i+1}, \dots, \tilde{s}_N).$$

At this time,

$$\begin{aligned} |\Phi(\mathcal{D}) - \Phi(\mathcal{D}^{(i)})| &= \left| \sup_{\boldsymbol{\theta} \in \Theta} (R(\boldsymbol{\theta}) - \widehat{R}_{\mathcal{D}}(\boldsymbol{\theta})) - \sup_{\boldsymbol{\theta} \in \Theta} (R(\boldsymbol{\theta}) - \widehat{R}_{\mathcal{D}^{(i)}}(\boldsymbol{\theta})) \right| \\ &\leq \sup_{\boldsymbol{\theta} \in \Theta} \left| (R(\boldsymbol{\theta}) - \widehat{R}_{\mathcal{D}}(\boldsymbol{\theta})) - (R(\boldsymbol{\theta}) - \widehat{R}_{\mathcal{D}^{(i)}}(\boldsymbol{\theta})) \right| \end{aligned}$$

holds. Here, in the last inequality, we used that generally

$$\left| \sup_{\boldsymbol{\theta}} a_{\boldsymbol{\theta}} - \sup_{\boldsymbol{\theta}} b_{\boldsymbol{\theta}} \right| \leq \sup_{\boldsymbol{\theta}} |a_{\boldsymbol{\theta}} - b_{\boldsymbol{\theta}}|$$

holds.

Furthermore, since $R(\boldsymbol{\theta}) = \mathbb{E}[g_{\boldsymbol{\theta}}(\tilde{s})]$ does not depend on the sample $\mathcal{D}$,

$$(R(\boldsymbol{\theta}) - \widehat{R}_{\mathcal{D}}(\boldsymbol{\theta})) - (R(\boldsymbol{\theta}) - \widehat{R}_{\mathcal{D}^{(i)}}(\boldsymbol{\theta})) = \widehat{R}_{\mathcal{D}^{(i)}}(\boldsymbol{\theta}) - \widehat{R}_{\mathcal{D}}(\boldsymbol{\theta})$$

holds. Therefore

$$|\Phi(\mathcal{D}) - \Phi(\mathcal{D}^{(i)})| \leq \sup_{\boldsymbol{\theta} \in \Theta} |\widehat{R}_{\mathcal{D}}(\boldsymbol{\theta}) - \widehat{R}_{\mathcal{D}^{(i)}}(\boldsymbol{\theta})|$$

is obtained.

Here, for arbitrary $\boldsymbol{\theta}$,

$$\widehat{R}_{\mathcal{D}}(\boldsymbol{\theta}) - \widehat{R}_{\mathcal{D}^{(i)}}(\boldsymbol{\theta}) = \frac{1}{N} (g_{\boldsymbol{\theta}}(\tilde{s}_i) - g_{\boldsymbol{\theta}}(\tilde{s}'_i)),$$

so from the assumption $0 \leq g_{\boldsymbol{\theta}}(\cdot) \leq B$,

$$|\widehat{R}_{\mathcal{D}}(\boldsymbol{\theta}) - \widehat{R}_{\mathcal{D}^{(i)}}(\boldsymbol{\theta})| \leq \frac{B}{N}$$

follows. Therefore

$$|\Phi(S) - \Phi(S^{(i)})| \leq \frac{B}{N}$$

is shown.

Therefore, from McDiarmid's inequality [McDiarmid, 1989], for arbitrary $t > 0$,

$$\Pr(\Phi(S) - \mathbb{E}[\Phi(S)] \geq t) \leq \exp\left(-\frac{2t^2}{\sum_{i=1}^N (B/N)^2}\right) = \exp\left(-\frac{2Nt^2}{B^2}\right)$$

holds. Therefore, for arbitrary $\eta_1 \in (0, 1)$, with probability at least $1 - \eta_1$,

$$\Phi(S) \leq \mathbb{E}[\Phi(S)] + B\sqrt{\frac{\log(1/\eta_1)}{2N}} \quad (47)$$

holds.

Next, regarding the empirical Rademacher complexity $\widehat{\mathfrak{R}}_{\mathcal{D}}(\mathcal{G})$, we also evaluate the amount of change due to replacing one sample point.

Let $\mathcal{D}^{(i)}$ be the sample obtained by replacing the i -th element of $\mathcal{D} = (\tilde{s}_1, \dots, \tilde{s}_N)$ with an independent identically distributed sample $\tilde{s}'_i$:

$$\mathcal{D}^{(i)} := (\tilde{s}_1, \dots, \tilde{s}_{i-1}, \tilde{s}'_i, \tilde{s}_{i+1}, \dots, \tilde{s}_N).$$

At this time, it suffices to evaluate

$$|\widehat{\mathfrak{R}}_{\mathcal{D}}(\mathcal{G}) - \widehat{\mathfrak{R}}_{\mathcal{D}^{(i)}}(\mathcal{G})| = \left| \mathbb{E}_{\sigma} \left[\sup_{g \in \mathcal{G}} \frac{1}{N} \sum_{j=1}^N \sigma_j g(\tilde{s}_j) \right] - \mathbb{E}_{\sigma} \left[\sup_{g \in \mathcal{G}} \frac{1}{N} \sum_{j=1}^N \sigma_j g(\tilde{s}_j^{(i)}) \right] \right|.$$

Here $\sigma = (\sigma_1, \dots, \sigma_N)$ are Rademacher variables.

Using the linearity of expectation and $|\mathbb{E}[X]| \leq \mathbb{E}[|X|]$,

$$|\widehat{\mathfrak{R}}_{\mathcal{D}}(\mathcal{G}) - \widehat{\mathfrak{R}}_{\mathcal{D}^{(i)}}(\mathcal{G})| \leq \mathbb{E}_{\sigma} \left[\left| \sup_{g \in \mathcal{G}} \frac{1}{N} \sum_{j=1}^N \sigma_j g(\tilde{s}_j) - \sup_{g \in \mathcal{G}} \frac{1}{N} \sum_{j=1}^N \sigma_j g(\tilde{s}_j^{(i)}) \right| \right].$$

Here, let

$$\psi_{\sigma}(\mathcal{D}) := \sup_{g \in \mathcal{G}} \frac{1}{N} \sum_{j=1}^N \sigma_j g(\tilde{s}_j),$$

then

$$|\widehat{\mathfrak{R}}_{\mathcal{D}}(\mathcal{G}) - \widehat{\mathfrak{R}}_{\mathcal{D}^{(i)}}(\mathcal{G})| \leq \mathbb{E}_{\sigma} \left[|\psi_{\sigma}(\mathcal{D}) - \psi_{\sigma}(\mathcal{D}^{(i)})| \right]$$

holds. Thus, we upper bound $|\psi_{\sigma}(\mathcal{D}) - \psi_{\sigma}(\mathcal{D}^{(i)})|$.

Since for any sequences of real numbers $\{a_g\}$ and $\{b_g\}$,

$$\left| \sup_g a_g - \sup_g b_g \right| \leq \sup_g |a_g - b_g|$$

holds, we have

$$\begin{aligned} |\psi_{\sigma}(\mathcal{D}) - \psi_{\sigma}(\mathcal{D}^{(i)})| &= \left| \sup_{g \in \mathcal{G}} \frac{1}{N} \sum_{j=1}^N \sigma_j g(\tilde{s}_j) - \sup_{g \in \mathcal{G}} \frac{1}{N} \sum_{j=1}^N \sigma_j g(\tilde{s}_j^{(i)}) \right| \\ &\leq \sup_{g \in \mathcal{G}} \left| \frac{1}{N} \sum_{j=1}^N \sigma_j g(\tilde{s}_j) - \frac{1}{N} \sum_{j=1}^N \sigma_j g(\tilde{s}_j^{(i)}) \right|. \end{aligned}$$

Here, since the difference between $\mathcal{D}$ and $\mathcal{D}^{(i)}$ is only in the i -th element, for any $g \in \mathcal{G}$,

$$\frac{1}{N} \sum_{j=1}^N \sigma_j g(\tilde{s}_j) - \frac{1}{N} \sum_{j=1}^N \sigma_j g(\tilde{s}_j^{(i)}) = \frac{1}{N} \sigma_i (g(\tilde{s}_i) - g(\tilde{s}'_i))$$

holds. From the assumption $0 \leq g(\cdot) \leq B$,

$$\left| \frac{1}{N} \sigma_i (g(\tilde{s}_i) - g(\tilde{s}'_i)) \right| \leq \frac{1}{N} |g(\tilde{s}_i) - g(\tilde{s}'_i)| \leq \frac{B}{N}.$$

Therefore,

$$|\psi_\sigma(\mathcal{D}) - \psi_\sigma(\mathcal{D}^{(i)})| \leq \frac{B}{N}$$

follows. Thus, we obtain

$$|\widehat{\mathfrak{R}}_{\mathcal{D}}(\mathcal{G}) - \widehat{\mathfrak{R}}_{\mathcal{D}^{(i)}}(\mathcal{G})| \leq \mathbb{E}_\sigma \left[|\psi_\sigma(\mathcal{D}) - \psi_\sigma(\mathcal{D}^{(i)})| \right] \leq \frac{B}{N}.$$

Therefore, by McDiarmid's inequality, for any $t > 0$,

$$\Pr \left(\widehat{\mathfrak{R}}_{\mathcal{D}}(\mathcal{G}) - \mathbb{E}_{\mathcal{D}}[\widehat{\mathfrak{R}}_{\mathcal{D}}(\mathcal{G})] \geq t \right) \leq \exp \left(- \frac{2Nt^2}{B^2} \right)$$

holds. Thus, for any $\eta_2 \in (0, 1)$, with probability at least $1 - \eta_2$,

$$\mathbb{E}_{\mathcal{D}}[\widehat{\mathfrak{R}}_{\mathcal{D}}(\mathcal{G})] \leq \widehat{\mathfrak{R}}_{\mathcal{D}}(\mathcal{G}) + B\sqrt{\frac{\log(1/\eta_2)}{2N}} \quad (48)$$

holds. That is,

$$\mathfrak{R}_N(\mathcal{G}) \leq \widehat{\mathfrak{R}}_{\mathcal{D}}(\mathcal{G}) + B\sqrt{\frac{\log(1/\eta_2)}{2N}}$$

Finally, in (47), by using $\mathbb{E}[\Phi(\mathcal{D})] \leq 2\mathfrak{R}_N(\mathcal{G})$ and further using (48) to upper bound $\mathfrak{R}_N(\mathcal{G})$ by $\widehat{\mathfrak{R}}_{\mathcal{D}}(\mathcal{G})$, specifically, evaluating the right-hand side of (47) as

$$\mathbb{E}[\Phi(\mathcal{D})] \leq 2\mathfrak{R}_N(\mathcal{G}) \leq 2\widehat{\mathfrak{R}}_{\mathcal{D}}(\mathcal{G}) + 2B\sqrt{\frac{\log(1/\eta_2)}{2N}},$$

we obtain

$$\Phi(\mathcal{D}) \leq 2\widehat{\mathfrak{R}}_{\mathcal{D}}(\mathcal{G}) + 2B\sqrt{\frac{\log(1/\eta_2)}{2N}} + B\sqrt{\frac{\log(1/\eta_1)}{2N}}.$$

Here, setting $\eta_1 = \eta_2 = \delta/2$, we have

$$2B\sqrt{\frac{\log(1/\eta_2)}{2N}} + B\sqrt{\frac{\log(1/\eta_1)}{2N}} = 3B\sqrt{\frac{\log(2/\delta)}{2N}}.$$

In this case, the probability that (47) and (48) hold simultaneously is at least $1 - (\eta_1 + \eta_2) = 1 - \delta$. Thus, with probability at least $1 - \delta$,

$$\Phi(\mathcal{D}) \leq 2\widehat{\mathfrak{R}}_{\mathcal{D}}(\mathcal{G}) + 3B\sqrt{\frac{\log(2/\delta)}{2N}}$$

is obtained. This is (34). □

D.8 PROOF OF COROLLARY B.5

Proof. From Theorem B.4, for any $\boldsymbol{\theta}$,

$$R(\boldsymbol{\theta}) \leq \widehat{R}_{\mathcal{D}}(\boldsymbol{\theta}) + 2\widehat{\mathfrak{R}}_{\mathcal{D}}(\mathcal{G}) + 3B\sqrt{\frac{\log(2/\delta)}{2N}}.$$

Applying this to $\boldsymbol{\theta} = \hat{\boldsymbol{\theta}}$ proves the claim. □

D.9 PROOF OF LEMMA C.2

Proof. First, we show $\text{Reg}(\hat{\mathbf{c}}, \mathbf{c}) \geq 0$. Since $z(\mathbf{c}) = \min_{\mathbf{w} \in \mathcal{W}} \mathbf{c}^\top \mathbf{w}$, for any $\mathbf{w} \in \mathcal{W}$, $z(\mathbf{c}) \leq \mathbf{c}^\top \mathbf{w}$ holds. In particular, taking $\mathbf{w} = \mathbf{w}^*(\hat{\mathbf{c}}) \in \mathcal{W}$, we have

$$\text{Reg}(\hat{\mathbf{c}}, \mathbf{c}) = \mathbf{c}^\top \mathbf{w}^*(\hat{\mathbf{c}}) - z(\mathbf{c}) \geq \mathbf{c}^\top \mathbf{w}^*(\hat{\mathbf{c}}) - \mathbf{c}^\top \mathbf{w}^*(\hat{\mathbf{c}}) = 0.$$

Next, we show the upper bound of Reg . Similar to the proof of Lemma B.2, for any $\mathbf{c}$,

$$\text{Reg}(\hat{\mathbf{c}}, \mathbf{c}) \leq \|\mathbf{c}\| \text{diam}_*(\mathcal{W}).$$

From Assumption C.1, since $\|\mathbf{c}\| \leq \bar{R}_C$,

$$\text{Reg}(\hat{\mathbf{c}}, \mathbf{c}) \leq \bar{R}_C \text{diam}_*(\mathcal{W}).$$

Furthermore, by definition, $\text{Lip}(\mathbf{w}^*(f_\theta(\mathbf{x})))$ is

$$\text{Lip}(\mathbf{w}^*(f_\theta(\mathbf{x}))) = \sup_{\mathbf{w} \in \mathcal{W}} \|\mathbf{w}^*(f_\theta(\mathbf{x})) - \mathbf{w}\|_*$$

and

$$0 \leq \text{Lip}(\mathbf{w}^*(f_\theta(\mathbf{x}))) \leq \text{diam}_*(\mathcal{W})$$

holds. Therefore,

$$0 \leq \varepsilon \text{Lip}(\mathbf{w}^*(f_\theta(\mathbf{x}))) \leq \varepsilon \text{diam}_*(\mathcal{W}).$$

Summing these up,

$$0 \leq g_\theta(\mathbf{x}, \mathbf{c}) = \text{Reg}(f_\theta(\mathbf{x}), \mathbf{c}) + \varepsilon \text{Lip}(\mathbf{w}^*(f_\theta(\mathbf{x}))) \leq (\bar{R}_C + \varepsilon) \text{diam}_*(\mathcal{W}) = B_{\text{dro}}.$$

Thus the claim holds. $\square$

D.10 PROOF OF THEOREM C.3

Proof. The proof follows the same framework as the proof of Theorem B.4. The only differences are

$$\mathcal{G} \rightarrow \mathcal{G}_{\text{dro}}, \quad B \rightarrow B_{\text{dro}}.$$

By Lemma C.2, g_θ falls within $[0, B_{\text{dro}}]$ almost surely, so the change due to the replacement of one sample point is bounded by B_{dro}/N . This yields the generalization bound with high probability ([Bartlett and Mendelson, 2002, McDiarmid, 1989]). $\square$

D.11 PROOF OF COROLLARY C.4

Proof. From Theorem C.3, for any θ ,

$$R_{\text{dro}}(\theta) \leq \hat{R}_{\text{dro}, \mathcal{D}}(\theta) + 2 \hat{\mathfrak{R}}_{\mathcal{D}}(\mathcal{G}_{\text{dro}}) + 3B_{\text{dro}} \sqrt{\frac{\log(2/\delta)}{2N}}.$$

Applying this to $\theta = \hat{\theta}$ proves the claim. $\square$

E DETAILS OF OPTIMIZATION PROBLEMS AND ALGORITHMS

E.1 DETAILS OF OPTIMIZATION PROBLEMS

Specific formulations of the optimization problems used in the experiments are shown below.

E.1.1 Shortest Path (SP)

We find the shortest path from a start node $s \in V$ to a target node $t \in V$ on a directed graph $G = (V, A)$. For each edge $a = (u, v) \in A$, a cost c_a is given, and the variable w_a represents whether to use edge a . The SP in PyEPO is formulated as follows:

$$\min_{\mathbf{w}} \sum_{(u,v) \in A} c_{uv} w_{uv} \quad (49)$$

$$\text{s.t.} \quad \sum_{(v,u) \in A} w_{vu} - \sum_{(u,v) \in A} w_{uv} = \begin{cases} -1 & (u = s), \\ 1 & (u = t), \\ 0 & (u \in V \setminus \{s, t\}), \end{cases} \quad (50)$$

$$0 \leq w_{uv} \leq 1 \quad (\forall (u, v) \in A). \quad (51)$$

E.1.2 Traveling Salesman Problem (TSP)

Consider TSP on a vertex set $V = \{1, \dots, m\}$ (in our experiments, $m = 8$). We introduce an undirected edge set $E = \{\{i, j\} \mid 1 \leq i < j \leq m\}$, where the variable $w_{ij} \in \{0, 1\}$ indicates whether edge $\{i, j\}$ is included in the tour. We use the following formulation:

$$\min_{\mathbf{w}} \sum_{1 \leq i < j \leq m} c_{ij} w_{ij} \quad (52)$$

$$\text{s.t.} \quad \sum_{j \in V: j \neq i} w_{\min\{i,j\}, \max\{i,j\}} = 2 \quad (\forall i \in V), \quad (53)$$

$$\sum_{\substack{1 \leq i < j \leq m \\ i \in S, j \in S}} w_{ij} \leq |S| - 1 \quad (\forall S \subset V, 2 \leq |S| \leq m - 1), \quad (54)$$

$$w_{ij} \in \{0, 1\} \quad (\forall \{i, j\} \in E). \quad (55)$$

E.1.3 Knapsack Problem (KP)

Consider a multi-dimensional ($L = 2$ in our experiments) 0–1 knapsack problem for a set of items $\{1, \dots, n\}$ ($n = 16$). The variable $x_i \in \{0, 1\}$ indicates whether item i is selected. Given values c_i , resource consumptions $a_{\ell i}$, and capacities b_ℓ , KP is given by the following maximization problem:

$$\max_{\mathbf{x}} \sum_{i=1}^n c_i x_i \quad (56)$$

$$\text{s.t.} \quad \sum_{i=1}^n a_{\ell i} x_i \leq b_\ell \quad (\ell = 1, \dots, L), \quad (57)$$

$$x_i \in \{0, 1\} \quad (i = 1, \dots, n). \quad (58)$$

E.1.4 Formulation of Energy Scheduling Problem

In the energy scheduling problem, we optimize the demand (power consumption) allocation over $T = 24$ periods (hourly). For each time $t \in \{1, \dots, T\}$, let the decision variable $w_t \in \mathbb{R}$ be the demand at that time. Let $\mathbf{p}^{\text{lo}}, \mathbf{p}^{\text{sch}}, \mathbf{p}^{\text{up}} \in \mathbb{R}^T$ be the lower bound of demand, baseline demand, and upper bound of demand, respectively.

The coefficient vector of the objective function $\mathbf{c} \in \mathbb{R}^T$ represents the unit cost of demand (e.g., electricity price) at each time, assumed to be generated or predicted from covariates $\mathbf{x} \in \mathbb{R}^p$. The energy scheduling problem in PyEPO is formulated as:

$$\mathbf{w}^*(\mathbf{c}) \in \arg \min_{\mathbf{w} \in \mathcal{W}} \mathbf{c}^\top \mathbf{w}, \quad \mathcal{W} := \left\{ \mathbf{w} \in \mathbb{R}^T \mid \mathbf{p}^{\text{lo}} \leq \mathbf{w} \leq \mathbf{p}^{\text{up}}, \mathbf{1}^\top \mathbf{w} = \mathbf{1}^\top \mathbf{p}^{\text{sch}} \right\}. \quad (59)$$

Here, $\mathbf{1} \in \mathbb{R}^T$ is a vector of all ones. The constraint $\mathbf{1}^\top \mathbf{w} = \mathbf{1}^\top \mathbf{p}^{\text{sch}}$ ensures that the total demand over 24 hours matches the total baseline demand.

E.2 DETAILS OF LEARNING ALGORITHMS

In this section, we describe the details of the learning algorithms in our proposed method.

E.2.1 Learning Algorithms Robust to Observation Errors

In learning using the robust regret loss described in Section 3.1, when using SPO+ [Elmachtoub and Grigas, 2022], the gradient of the SPO+ loss with $\mathbf{c}^*$ regarded as the true coefficient vector,

$$G_{\text{SPO+}}(\hat{\mathbf{c}}; \mathbf{c}^*) := 2 \left(\mathbf{w}^*(\mathbf{c}^*) - \mathbf{w}^*(2\hat{\mathbf{c}} - \mathbf{c}^*) \right) \quad (60)$$

is used as an approximation of $\nabla_{\hat{\mathbf{c}}} \mathcal{L}_{\text{rob}}$. When using FY loss [Blondel et al., 2020], using the optimal solution with regularization term Ω , $\hat{\mathbf{w}}_{\Omega}(\hat{\mathbf{c}}) = \arg \max_{\mathbf{w} \in \mathcal{W}} \{ \langle \hat{\mathbf{c}}, \mathbf{w} \rangle - \Omega(\mathbf{w}) \}$,

$$G_{\text{FY}}(\hat{\mathbf{c}}; \mathbf{c}^*) := \hat{\mathbf{w}}_{\Omega}(\hat{\mathbf{c}}) - \mathbf{w}^*(\mathbf{c}^*) \quad (61)$$

is used as the gradient. Here $\mathbf{w}^*(\mathbf{c}^*)$ is the optimal solution for $\mathbf{c}^*$.

Summarizing the above procedure, the learning algorithm by the proposed method iterates the following.

1. Based on the current model output $\hat{\mathbf{c}}$, solve the maximization problem of the robust regret loss to find the worst-case coefficient vector $\mathbf{c}^*$.
2. For the obtained $\mathbf{c}^*$, compute the gradient $G(\hat{\mathbf{c}}; \mathbf{c}^*)$ of the surrogate loss with respect to $\hat{\mathbf{c}}$.
3. Update the model parameter $\boldsymbol{\theta}$ by backpropagation using this gradient G :

$$\boldsymbol{\theta} \leftarrow \boldsymbol{\theta} - \eta \cdot \left(\frac{\partial \hat{\mathbf{c}}}{\partial \boldsymbol{\theta}} \right)^{\top} G(\hat{\mathbf{c}}; \mathbf{c}^*)$$

Here, η is the learning rate. This allows learning a predictive model that makes robust decisions even for the worst case within the uncertainty set.

E.2.2 Learning Algorithms Robust to Distribution Shifts

The details of the learning algorithm using the distributionally robust regret loss described in Section 3.2 are as follows.

The learning algorithm using distributionally robust regret iterates the following.

1. Based on the current model output $\hat{\mathbf{c}} = f_{\boldsymbol{\theta}}(\mathbf{x})$, find the decision $\mathbf{w}^*(\hat{\mathbf{c}})$, and evaluate the distributionally robust loss

$$\mathcal{L}_{\text{dro}}(\mathbf{x}) = \mathbb{E}_{\tilde{\mathbf{c}} \sim P_{\tilde{\mathbf{c}}|\mathbf{x}}} [\text{Reg}(\hat{\mathbf{c}}, \tilde{\mathbf{c}})] + \varepsilon \text{Lip}(\mathbf{w}^*(\hat{\mathbf{c}}))$$

Here $\text{Lip}(\mathbf{w}^*(\hat{\mathbf{c}})) = \sup_{\mathbf{w} \in \mathcal{W}} \|\mathbf{w}^*(\hat{\mathbf{c}}) - \mathbf{w}\|_*$.

2. Compute surrogate gradients for the two terms in the above equation. Specifically, for the regret term, compute $G_{\text{reg}}(\hat{\mathbf{c}}; \tilde{\mathbf{c}})$ using existing surrogate gradients such as SPO+ or FY loss. For the Lipschitz term, using the optimal solution

$$(\mathbf{w}^*, \mathbf{g}^*) \in \arg \max_{(\mathbf{w}, \mathbf{g}) \in \mathcal{W} \times \mathcal{G}} \mathbf{g}^{\top} (\mathbf{w}^*(\hat{\mathbf{c}}) - \mathbf{w}),$$

compute its subgradient (or approximate gradient) $G_{\text{lip}}(\hat{\mathbf{c}}; \mathbf{g}^*)$. Combining these, we use

$$G_{\text{dro}}(\hat{\mathbf{c}}) := G_{\text{reg}}(\hat{\mathbf{c}}) + \varepsilon G_{\text{lip}}(\hat{\mathbf{c}})$$

as the approximate gradient of the distributionally robust loss.

3. Update the model parameter $\boldsymbol{\theta}$ by backpropagation using the obtained gradient $G_{\text{dro}}(\hat{\mathbf{c}})$:

$$\boldsymbol{\theta} \leftarrow \boldsymbol{\theta} - \eta \cdot \left(\frac{\partial \hat{\mathbf{c}}}{\partial \boldsymbol{\theta}} \right)^{\top} G_{\text{dro}}(\hat{\mathbf{c}}).$$

Here, η is the learning rate. This allows learning a predictive model that suppresses the deterioration of regret even when the conditional distribution of the coefficient vector is perturbed in the sense of Wasserstein distance.

F DETAILS OF EXPERIMENTAL SETTINGS

F.1 COMMON SETTINGS

Common Settings for Datasets and Learning In our experiments, the following common settings were used for all optimization problems.

- **Data Splitting:** The total number of samples N_{total} was split into training (N_{train}), validation (N_{val}), and test (N_{test}) sets. Specifically, we set $N_{\text{train}} = 100$, $N_{\text{val}} = 100$, and $N_{\text{test}} = 1000$.
- **Dimension of Covariates:** $p = 5$.
- **Polynomial Degree:** A polynomial model of degree $d = 4$ was used for data generation (see Section F.2 for details).
- **Learning Settings:** Adam [Kingma and Ba, 2015] was used as the optimization algorithm, with a learning rate of 10^{-2} , ℓ_2 regularization coefficient of 10^{-4} , batch size of 64, and 200 epochs.
- **Size of Optimization Problems:** For SP, we used a 5×5 grid graph; for TSP, the number of vertices $m = 8$; for KP, the number of items $n = 16$ and the number of capacity constraints $L = 2$.

Settings for Perturbation Parameters **Parameters for Observation Error (Parameter Noise):** For sparse perturbation (Sparse), the ratio of selected dimensions was $\gamma_{\text{frac}} = 0.5$. The shift sizes δ were as follows:

- SP: $\delta = 5.0$
- TSP: $\delta = 28.0$
- KP: $\delta = 22.0$

Parameters for Distribution Shift (Distribution Noise): The parameters for each problem setting are as follows.

- SP: $d_\gamma = 4$, $\alpha_\gamma = 0.8$, $\sigma_{\text{global}} = 1.0$.
- TSP: $d_\gamma = 4$, $\alpha_\gamma = 1.0$, $\sigma_{\text{global}} = 0.30$.
- KP: $d_\gamma = 4$, $\alpha_\gamma = 0.6$, $\sigma_{\text{global}} = 0.50$.

Settings for Real Data Experiments In the experiments using real data, the total number of data points was 362, consisting of 100 training samples, 100 validation samples, and 162 test samples.

Details of Evaluation Metrics To assess robustness against worst-case scenarios, we report the Conditional Value-at-Risk (CVaR_α) at confidence level $\alpha = 0.9$. Specifically, let $\tilde{r}_{(1)} \leq \dots \leq \tilde{r}_{(N_{\text{test}})}$ denote the sorted relative regrets on the test set. We compute CVaR_α as the average of the worst $(1 - \alpha)$ -fraction of regrets:

$$\text{CVaR}_\alpha(\mathbf{R}) := \frac{1}{N_{\text{test}} - k_\alpha + 1} \sum_{j=k_\alpha}^{N_{\text{test}}} \tilde{r}_{(j)}, \quad (62)$$

where $k_\alpha := \lceil \alpha N_{\text{test}} \rceil$.

F.2 DETAILS OF SYNTHETIC DATA GENERATION

For an n -dimensional coefficient vector, we fix a 0–1 random matrix $B \in \{0, 1\}^{n \times p}$, and generate the true coefficient vector $\mathbf{c}_i \in \mathbb{R}^n$ calculated from the covariate $\mathbf{x}_i \sim \mathcal{N}(\mathbf{0}, I_p)$ by the following procedure. Here, $\mathcal{N}(\mathbf{0}, I_p)$ denotes the multivariate normal distribution with mean vector $\mathbf{0}$ and covariance matrix I_p . First, we compute the base vector $\tilde{\mathbf{c}}_i$ as

$$\tilde{\mathbf{c}}_i = \left(\frac{B\mathbf{x}_i}{\sqrt{p}} + 3\mathbf{1} \right)^d + \mathbf{1} \quad (63)$$

and set the final coefficient vector $\mathbf{c}_i$ as

$$\mathbf{c}_i = \frac{1}{3.5^d} \tilde{\mathbf{c}}_i \circ \boldsymbol{\epsilon}_i. \quad (64)$$

Here, $d = 4$ is the degree of the polynomial, $\circ$ denotes element-wise product, and $\mathbf{1}$ is a vector of all ones. Also, $\boldsymbol{\epsilon}_i \in \mathbb{R}^n$ is a multiplicative noise term, where each element follows a uniform distribution $\text{Unif}[1 - \eta, 1 + \eta]$.

F.3 DETAILS OF DISTRIBUTION SHIFT GENERATION

We describe the detailed procedure for generating distribution shifts. Starting from coefficient vectors $\mathbf{c}_i$ generated from the noiseless distribution, we add noise according to a mixture perturbation model that combines (i) a covariate-dependent bias term and (ii) covariate-independent random noise. For the covariate-dependent perturbation, we define a polynomial map $g_\gamma : \mathbb{R}^p \rightarrow \mathbb{R}^n$ using a binary random matrix $B \in \{0, 1\}^{n \times p}$ and degree $d_\gamma \in \mathbb{N}$, and set

$$\mathbf{c}_{i,\text{obs}} = \mathbf{c}_i + \alpha_\gamma g_\gamma(\mathbf{x}_i).$$

For the covariate-independent perturbation, we sample

$$\boldsymbol{\varepsilon}_i \sim \mathcal{N}(\mathbf{0}, \sigma_{\text{global}}^2 I_n)$$

and set

$$\mathbf{c}_{i,\text{obs}} = \mathbf{c}_i + \boldsymbol{\varepsilon}_i.$$

We switch between these two perturbations with probability π_{mix} to obtain the final observed coefficient vector. The specific parameter values used in the experiments (e.g., degree d_γ , coefficient α_γ , noise variance σ_{global}) are listed in Section F.1.

F.4 HYPERPARAMETER CANDIDATES

In the hyperparameter search performed using the validation dataset in our experiments, the following parameter sets were considered as candidates.

- SP:

$$\kappa \in \{2.0, 7.5, 12.5, 20.0, 30.0\}, \quad (65)$$

$$\rho \in \{1.0, 3.0, 5.0, 7.0, 9.0\}. \quad (66)$$

- KP:

$$\kappa \in \{1.5, 4.0, 9.0, 15.0, 25.0\}, \quad (67)$$

$$\rho \in \{10.0, 14.0, 22.0, 28.0, 48.0\}. \quad (68)$$

- TSP:

$$\kappa \in \{1.0, 3.0, 8.0, 18.0, 30.0, 50.0\}, \quad (69)$$

$$\rho \in \{10.0, 24.0, 28.0, 35.0, 42.0, 65.0\}. \quad (70)$$

In the experiments on robustness against observation errors, the parameter of the budget uncertainty set $\mathcal{C}_\Gamma$ in our proposed method was set to $\Gamma = n/2$, and the identity matrix was used for the covariance matrix Σ of the ellipsoidal uncertainty set $\mathcal{C}_{\text{ell}}$. The size parameters ρ and κ of the uncertainty sets were selected from the candidate sets above using the validation dataset.

In the experiments on robustness against distribution shifts, the Wasserstein ball radius ε was selected from the set $\{0.25, 0.5, 1.0, 2.0\}$ (for KP only $\{0.05, 0.1, 0.25, 0.5, 1.0, 2.0\}$), and determined based on regret on the validation data.

F.5 COMPUTATIONAL ENVIRONMENT

All numerical experiments in this study were conducted under the following computational environment.

- OS: macOS 15.3
- Processor: Apple M4
- RAM: 16GB
- Language: Python 3.9.6
- Optimization Software: Gurobi 12.0.3

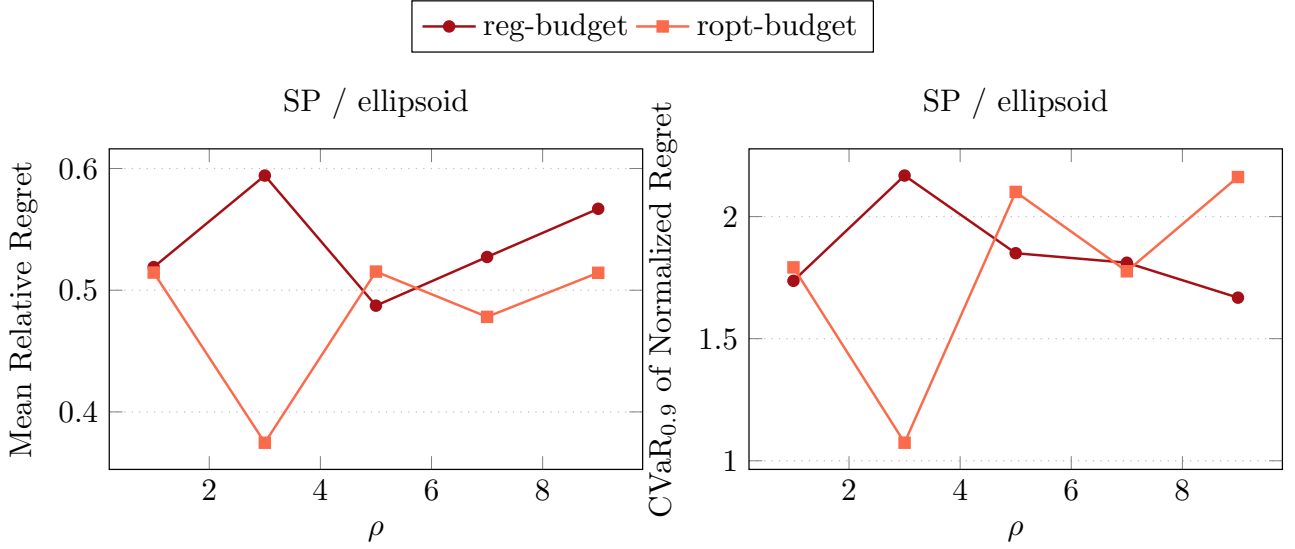


Figure 8: Sensitivity of mean relative regret and CVaR_{0.9} for ropt-budget in SP (under ellipsoidal perturbation) (Parameter ρ).

G ADDITIONAL EXPERIMENTAL RESULTS

G.1 SENSITIVITY ANALYSIS

In this section, we describe the results of sensitivity analysis regarding the size parameters (ρ, κ) of the uncertainty sets. As shown in Figure 8, for SP (under ellipsoidal perturbation), a trade-off can be confirmed: robustness is insufficient when ρ is too small, while excessive conservativeness worsens regret when ρ is too large. As a result, CVaR_{0.9} is minimized around $\rho = 3$.

Similarly, as can be seen from Figure 9, in TSP as well, regret worsens if κ is either too small or too large, and both mean relative regret and CVaR_{0.9} are best around $\kappa = 18$. Note that since ρ and κ are parameters corresponding to uncertainty sets of different shapes, simple comparison on the same numerical scale is not possible. Therefore, appropriate candidate sets must be set for each problem setting and selected based on validation data. From the above results, it was shown that the proposed method can improve both mean relative regret and CVaR_{0.9} in situations where observation errors exist. On the other hand, since the improvement effect depends on the choice of the uncertainty set and the setting of hyperparameters, an appropriate verification process according to the scale of errors assumed during operation is important.

Figure 10 presents a sensitivity analysis of the mean relative regret in KP. We observe a clear trade-off in the Wasserstein radius ε : when ε is too small, the method behaves similarly to emp, whereas an excessively large ε leads to overly conservative decisions. The choice of ground norm is also important; for example, the uncertainty set can expand rapidly under ℓ_∞ , suggesting that the norm and radius should be selected to match the problem characteristics. Overall, these results confirm that distributionally robust DFL with appropriately chosen norm and radius is effective under distribution shift.

Furthermore, Figure 11 shows the results of sensitivity analysis of CVaR_{0.9} with respect to the Wasserstein radius and norm in KP. Similar to the mean relative regret, although a tendency to become overly conservative is seen when the radius ε is excessive, it was confirmed that CVaR_{0.9} can be effectively reduced by selecting an appropriate ε . Also, the ℓ_∞ norm has higher sensitivity to parameters compared to other norms, suggesting that appropriate parameter setting is more important.

H RELATED WORK

In this section, we review previous research related to this study. First, in Section H.1, we describe two-stage methods that separate learning and optimization. Next, in Section H.2, we explain robust optimization and distributionally robust optimization considering the uncertainty of coefficient vectors. Then, in Section H.3, we overview the research trends of Decision-Focused Learning (DFL), and finally in Section H.4, we summarize existing research on robust DFL closely

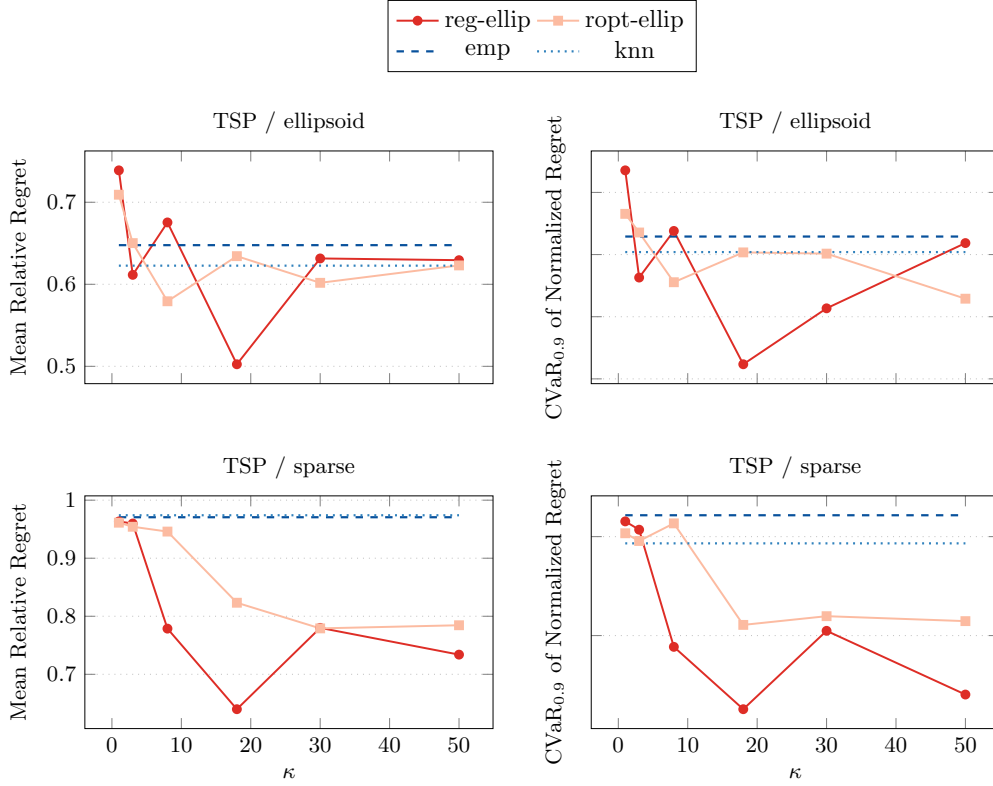


Figure 9: Sensitivity of mean relative regret and $\text{CVaR}_{0.9}$ of the loss using the ellipsoidal uncertainty set (`reg-ellip`) in TSP (Parameter κ).

related to this study.

H.1 TWO-STAGE METHODS SEPARATING LEARNING AND OPTIMIZATION

When the coefficient vector of the objective function in a mathematical optimization problem is unknown, its value needs to be predicted using observable covariates. Predict-then-Optimize takes a two-stage procedure: first predicting the coefficient vector from covariates using a predictive model, and then solving the subsequent optimization problem using the obtained predicted value. Examples include demand prediction and order quantity determination in the newsvendor problem [Huber et al., 2019, Rios et al., 2015] and prescriptive analysis based on conditional expectation estimation [Bertsimas and Kallus, 2020], which are major methods in contextual optimization [Sadana et al., 2025]. However, it has been pointed out that the minimization of the loss function (such as mean squared error) in the first learning stage does not necessarily coincide with the improvement of decision quality (i.e., minimization of regret) in the subsequent optimization problem [Sadana et al., 2025].

H.2 OPTIMIZATION PROBLEMS CONSIDERING UNCERTAINTY OF COEFFICIENT VECTORS

Robust optimization and distributionally robust optimization are known as methods to directly model the uncertainty of the coefficient vector and suppress performance degradation in the worst case. In robust optimization, it is assumed that the coefficient vector takes a value within a certain uncertainty set, and optimization is performed assuming the worst case within that set [Ben-Tal et al., 2009, Bertsimas et al., 2011, Bertsimas and Sim, 2004]. On the other hand, DRO assumes that the true distribution of the coefficient vector is unknown, and minimizes the worst-case expected value over an uncertainty set of distributions (such as sets based on moment conditions or Wasserstein distance) [Delage and Ye, 2010, Wiesemann et al., 2014, Mohajerin Esfahani and Kuhn, 2018]. Although these are usually optimization frameworks that do not include the learning process of predictive models, they provide fundamental ideas for dealing with uncertainty.

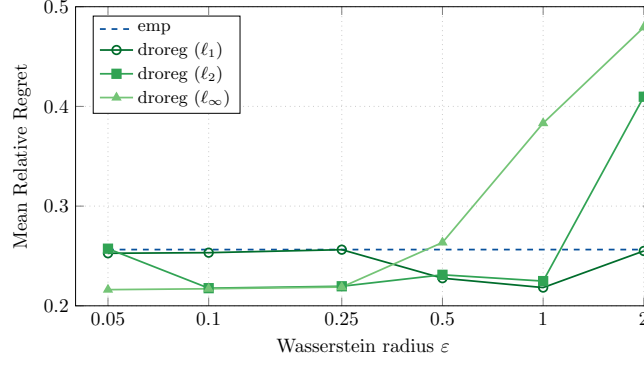


Figure 10: Sensitivity of mean relative regret to Wasserstein radius and norm in KP.

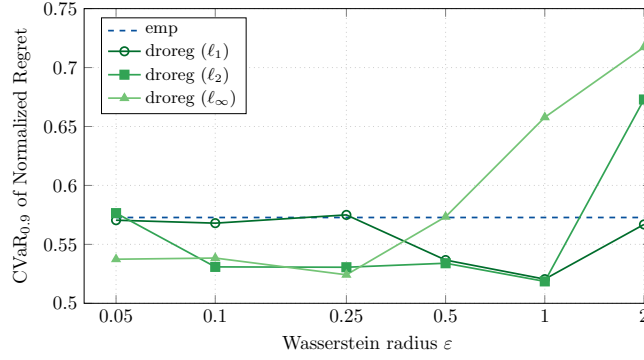


Figure 11: Sensitivity of $\text{CVaR}_{0.9}$ to Wasserstein radius and norm in KP.

H.3 DECISION-FOCUSED LEARNING

To address the mismatch between learning and optimization mentioned above, Decision-Focused Learning (DFL) has been proposed, which learns a model to directly improve the quality of the final decision (e.g., regret). DFL methods are broadly classified into two categories.

H.3.1 Methods Based on Implicit Differentiation and Continuous Relaxation

One is to treat the optimal solution of the optimization problem as a differentiable layer in a neural network and compute gradients using implicit differentiation or continuous relaxation. These include OptNet which incorporates quadratic programming as a layer [Amos and Kolter, 2017], CVXPY Layers providing convex optimization layers [Agrawal et al., 2019], and DDP simplifying dynamic programming [Mensch and Blondel, 2018]. Methods using continuous relaxation or cutting plane methods have also been proposed for MIP (Mixed Integer Programming) [Wilder et al., 2019, Ferber et al., 2020, Donti et al., 2017].

H.3.2 Methods Based on Surrogate Loss and Surrogate Gradient

The second approach is to design surrogate loss functions or surrogate gradients for regret to deal with the discontinuity of the optimal solution or when the solver is a black box. Representative methods include SPO+ loss giving a convex upper bound of regret [Elmachtoub and Grigas, 2022] and Fenchel-Young loss based on Fenchel duality [Blondel et al., 2020]. Perturbation Gradient methods approximating directional derivatives with finite differences [Huang and Gupta, 2024], methods learning the shape of the loss function [Shah et al., 2022], and gradient estimation for black-box solvers [Vlastelica et al., 2020] have also been proposed.

H.4 ROBUST DECISION-FOCUSED LEARNING

Recently, research has been progressing on enhancing the robustness of DFL against data and environmental uncertainty. For example, Schutte et al. [2024] proposed RO loss based on the concept of robust optimization, risk-averse top- k loss, and k -NN loss which performs smoothing using neighboring data to address observation errors and data scarcity. Ma et al. [2024] proposed a method to incorporate a layer solving the DRO problem into a neural network. Furthermore, Wang et al. [2026] proposed Gen-DFL, which learns the conditional distribution of coefficient vectors using a generative model and minimizes CVaR on the learned distribution.