

Online Supplement of “Decision-Driven Regularization: A Blended Model for Learning and Optimization”

EC1. Omitted Proofs

In this segment, we present all deferred proofs from the main text.

EC1.1. Proof of Claim 1 in Illustration 1

The least squares minimization problem for OLS has the solution: $\mathbf{W} = (\mathbf{X}^\top \mathbf{X})^{-1} (\mathbf{X}^\top \mathbf{Z})$, where $\mathbf{X} = \begin{bmatrix} \tilde{\mathbf{x}}_1^\top & \tilde{\mathbf{x}}_2^\top & \tilde{\mathbf{x}}_3^\top & \tilde{\mathbf{x}}_4^\top \end{bmatrix}^\top$ and $\mathbf{Z} = (\tilde{z}_1^\top, \tilde{z}_2^\top, \tilde{z}_3^\top, \tilde{z}_4^\top)^\top$. It is easy to verify that these equations return $\hat{\mathbf{W}}^{\text{SLO}} = \check{\mathbf{W}} = \begin{bmatrix} 1 & 0 \\ 0 & 1.5 \end{bmatrix}$. Hence, SLO picks route b if and only if $1.5\hat{x}_2 \leq \hat{x}_1$.

To verify the conclusion for ILO, note that for any given feature vector $\hat{\mathbf{x}}$ and coefficient matrix $\mathbf{W}$, route b is chosen if and only if $(\mathbf{W}\hat{\mathbf{x}})_1 - (\mathbf{W}\hat{\mathbf{x}})_2 = w_{11}\hat{x}_1 + w_{12}\hat{x}_2 - w_{21}\hat{x}_1 - w_{22}\hat{x}_2 \geq 0$. For convenience, let $\bar{w} := w_{11} - w_{21}$ and $\underline{w} := w_{22} - w_{12}$. Hence, the optimal decision $\mathbf{y}^*$ will pick route b if and only if $\underline{w}\hat{x}_2 \leq \bar{w}\hat{x}_1$.

For ILO, in order to get the minimum objective, ILO chooses $\hat{\mathbf{W}}^{\text{ILO}}$ such that the smaller of the two components of $\tilde{\mathbf{z}}_n$ is chosen. Specifically, to minimize ILO loss, we need to pick route b for data points 1, 3 and 4, but pick route a for data point 2. Hence, optimal $\hat{\mathbf{W}}^{\text{ILO}}$ is attained when

$$\begin{aligned} \bar{w} - \underline{w} &\geq 0, \\ 2\bar{w} - \underline{w} &\geq 0, \\ \bar{w} - 2\underline{w} &\leq 0, \\ 3\bar{w} - 2\underline{w} &\geq 0. \end{aligned}$$

Subtracting the LHS of the first inequality from the third one, we have $\bar{w} - \underline{w} - (\bar{w} - 2\underline{w}) \geq 0$, i.e., $\underline{w} \geq 0$. Then the above four inequalities are equivalent to the condition $1 \leq \bar{w}/\underline{w} \leq 2$. In other words, any matrix $\hat{\mathbf{W}}^{\text{ILO}}$ satisfying $1 \leq \bar{w}/\underline{w} \leq 2$, $\underline{w} \geq 0$ is optimal for the ILO loss function. Thus, ILO chooses route b if and only if $\hat{x}_2 \leq (\bar{w}/\underline{w})\hat{x}_1$ for some $1 \leq \bar{w}/\underline{w} \leq 2$ that is degenerate and indifferentiable in-sample under the ILO loss. $\square$

EC1.2. Proof of Proposition 1

$$\begin{aligned} &\mathbb{E}_{\mathcal{D}^N} [(\mathbf{y}^\top \mathbb{E}_{\mathbf{z}|\mathbf{x}}[\mathbf{z}] - \mathbf{y}^\top \mathbf{f}(\mathbf{x}; \hat{\mathbf{w}}))^2] \\ &= \mathbb{E}_{\mathcal{D}^N} [(\mathbf{y}^\top g(\mathbf{x}) - \mathbf{y}^\top \mathbf{f}(\mathbf{x}; \hat{\mathbf{w}}))^2] \\ &= \mathbb{E}_{\mathcal{D}^N} [\mathbf{y}^\top g(\mathbf{x})]^2 - 2\mathbb{E}_{\mathcal{D}^N} [\mathbf{y}^\top g(\mathbf{x}) \mathbf{y}^\top \mathbf{f}(\mathbf{x}; \hat{\mathbf{w}})] + \mathbb{E}_{\mathcal{D}^N} [(\mathbf{y}^\top \mathbf{f}(\mathbf{x}; \hat{\mathbf{w}}))^2] \\ &= [\mathbf{y}^\top g(\mathbf{x})]^2 - 2\mathbf{y}^\top g(\mathbf{x}) \mathbb{E}_{\mathcal{D}^N} [\mathbf{y}^\top \mathbf{f}(\mathbf{x}; \hat{\mathbf{w}})] + \mathbb{E}_{\mathcal{D}^N} [(\mathbf{y}^\top \mathbf{f}(\mathbf{x}; \hat{\mathbf{w}}))^2] - (\mathbb{E}_{\mathcal{D}^N} [\mathbf{y}^\top \mathbf{f}(\mathbf{x}; \hat{\mathbf{w}})])^2 + (\mathbb{E}_{\mathcal{D}^N} [\mathbf{y}^\top \mathbf{f}(\mathbf{x}; \hat{\mathbf{w}})])^2 \\ &= [\mathbf{y}^\top g(\mathbf{x}) - \mathbb{E}_{\mathcal{D}^N} \mathbf{y}^\top \mathbf{f}(\mathbf{x}; \hat{\mathbf{w}})]^2 + \mathbb{E}_{\mathcal{D}^N} [(\mathbf{y}^\top \mathbf{f}(\mathbf{x}; \hat{\mathbf{w}}))^2] - (\mathbb{E}_{\mathcal{D}^N} [\mathbf{y}^\top \mathbf{f}(\mathbf{x}; \hat{\mathbf{w}})])^2. \end{aligned}$$

where the first equality holds due to $\mathbb{E}_{\mathbf{z}|\mathbf{x}}[\mathbf{z}] = g(\mathbf{x})$, the third equality holds because $\mathbf{y}$ and $\mathbf{x}$ are deterministic with respect to the data $\mathcal{D}^N$ and the last equality comes from rearranging terms.

EC1.3. Proof of Proposition 2

Let (α_1, α_2) be the dual variables, then the Lagrangian relaxation of Problem (4) is given by:

$$\begin{aligned} & \min_{\mathbf{y} \in \mathcal{Y}} \left\{ \mathbf{y}^\top \mathbf{f}(\tilde{\mathbf{x}}, \mathbf{w}) + \alpha_1 \left(\mathbf{y}^\top \tilde{\mathbf{z}} - \mathbf{y}^\top \mathbf{f}(\tilde{\mathbf{x}}, \mathbf{w}) - \gamma \right) + \alpha_2 \left(-\mathbf{y}^\top \tilde{\mathbf{z}} + \mathbf{y}^\top \mathbf{f}(\tilde{\mathbf{x}}, \mathbf{w}) - \gamma \right) \right\} \\ & = \min_{\mathbf{y} \in \mathcal{Y}} \left\{ (\alpha_1 - \alpha_2) \mathbf{y}^\top \tilde{\mathbf{z}} + (1 - \alpha_1 + \alpha_2) \mathbf{y}^\top \mathbf{f}(\tilde{\mathbf{x}}, \mathbf{w}) - \gamma(\alpha_1 + \alpha_2) \right\}. \end{aligned}$$

Setting $\mu = \alpha_1 - \alpha_2$, we recover the statement of the Proposition 2, where the term associated with γ is dropped as it do not influence the choice of decisions $\mathbf{y}$.

EC1.4. Proof of Proposition 3

When $c(\mathbf{y}; \mathbf{z}) = \mathbf{y}^\top \mathbf{z}$ and $\mathcal{Y} = \{\mathbf{A}\mathbf{y} \leq \mathbf{b}, \mathbf{y} \geq \mathbf{0}\}$, the inner supreme problem of the DDR becomes

$$\begin{aligned} & \max_{\mathbf{y}_n \geq \mathbf{0}} \mathbf{y}_n^\top \left(-\mu \tilde{\mathbf{z}}_n - (1 - \mu) \mathbf{f}(\tilde{\mathbf{x}}_n; \mathbf{w}) \right) \\ & \text{s.t. } \mathbf{A}\mathbf{y}_n \leq \mathbf{b}, \end{aligned}$$

which is a linear programming problem and the dual of it is

$$\begin{aligned} & \min_{\boldsymbol{\kappa}_n \geq \mathbf{0}} \boldsymbol{\kappa}_n^\top \mathbf{b} \\ & \text{s.t. } \mathbf{A}^\top \boldsymbol{\kappa}_n \geq -\mu \tilde{\mathbf{z}}_n - (1 - \mu) \mathbf{f}(\tilde{\mathbf{x}}_n; \mathbf{w}). \end{aligned}$$

□

EC1.5. Motivation for the loss function uncertainty set

The following illustration gives some explanation about why this geometry for the uncertainty set makes sense.

Illustration 1 [as seen in [Zhu et al. \(2022\)](#)]. Suppose $L(\mathbf{w}) = -2 \log \left(\prod_n p_{\mathbf{z}|\mathbf{x}}(\mathbf{z}_n; \mathbf{x}_n, \mathbf{w}) \right)$ was chosen as the log-likelihood function, where $p_{\mathbf{z}|\mathbf{x}}$ is the density of $\mathbf{z}|\mathbf{x}$. Then the uncertainty set $\mathcal{U}(\rho) := \{\mathbf{w} : L(\mathbf{w}) - L(\tilde{\mathbf{w}}) \leq \rho\}$ reduces to the set of all weights $\mathbf{w}$, which likelihood ratio against $\tilde{\mathbf{w}}$ is no greater than some bound:

$$\mathcal{U}(\rho) := \left\{ \mathbf{w} : LR(\tilde{\mathbf{w}}; \mathbf{w}) := \frac{\prod_n p_{\mathbf{z}|\mathbf{x}}(\mathbf{z}_n; \mathbf{x}_n, \tilde{\mathbf{w}})}{\prod_n p_{\mathbf{z}|\mathbf{x}}(\mathbf{z}_n; \mathbf{x}_n, \mathbf{w})} \leq e^{\rho/2} \right\}. \quad (\text{EC1.1})$$

Suppose we interpret the null hypothesis as $\mathcal{H}_0 : \tilde{\mathbf{w}} = \tilde{\mathbf{w}}$ and the alternative hypothesis as $\mathcal{H}_1 : \tilde{\mathbf{w}} = \mathbf{w} \neq \tilde{\mathbf{w}}$. Then Neyman-Pearson Lemma gives that there is some confidence level $\alpha(\rho)$, corresponding to $e^{\rho/2}$, under which, the likelihood ratio test grants the highest power, that is, probability of rejecting $\mathcal{H}_0$. In other words, this uncertainty set is equivalent to considering the maximal set of all weights $\mathbf{w}$ that, if true, have any chance at all of rejecting $\tilde{\mathbf{w}}$ under the significance level corresponding to $e^{\rho/2}$, given the existing dataset.

EC1.6. Proof of Theorem 1

Before proving Theorem 1, we first prove the following lemma.

LEMMA 1. *The valuation function $v_\mu(\mathbf{w})$ is concave in $\mathbf{w}$.*

Proof. Under the assumption that $\mathbf{f}(\mathbf{x}; \mathbf{w})$ is concave in $\mathbf{w}$ for all $\mathbf{x} \in \mathcal{X}$, we see that $\mathbf{y}^\top \mathbf{f}(\mathbf{x}, \mathbf{w})$ is convex in $\mathbf{y}$ and concave in $\mathbf{w}$ for all $\mathbf{x}$. It follows that $\mathbf{y}^\top \tilde{\mathbf{z}}$ is convex in $\mathbf{y}$. Now for any $\mu \in [0, 1]$, it is easy to see that the objective function $\mu \mathbf{y}^\top \tilde{\mathbf{z}} + (1 - \mu) \mathbf{y}^\top \mathbf{f}(\mathbf{x}; \mathbf{w})$ is convex in $\mathbf{y}$ and concave in $\mathbf{w}$. The concavity of $v_\mu(\mathbf{w})$ over $\mathbf{w}$ arises from the infimum operator over $\mathbf{y}$. $\square$

Now we are ready to prove Theorem 1. For **DDR**, we have

$$\begin{aligned} & \min_{\mathbf{w}} \left\{ L(\mathbf{w}) - \lambda \frac{1}{N} \sum_{n \in [N]} \min_{\mathbf{y}_n \in \mathcal{Y}} \left\{ \mu \mathbf{y}_n^\top \tilde{\mathbf{z}}_n + (1 - \mu) \mathbf{y}_n^\top \mathbf{f}(\tilde{\mathbf{x}}_n; \mathbf{w}) \right\} \right\} \\ &= \min_{\mathbf{w}} \left\{ \frac{1}{N} \sum_{n \in [N]} \max_{\mathbf{y}_n \in \mathcal{Y}} \left\{ \ell(\tilde{\mathbf{z}}_n; \mathbf{f}(\tilde{\mathbf{x}}_n; \mathbf{w})) - \lambda \left(\mu \mathbf{y}_n^\top \tilde{\mathbf{z}}_n + (1 - \mu) \mathbf{y}_n^\top \mathbf{f}(\tilde{\mathbf{x}}_n; \mathbf{w}) \right) \right\} \right\}. \end{aligned} \quad (\text{EC1.2})$$

In what follows, we denote $\sigma = L(\tilde{\mathbf{w}})$, which is a known constant given the data. For **Robust-DDR**, we can write the Lagrangian of it as

$$\begin{aligned} & \max_{\mathbf{y}_n \in \mathcal{Y}, n \in [N], \alpha \geq 0} \min_{\mathbf{w}} \left\{ \alpha \left(L(\mathbf{w}) - \sigma - \rho \right) - \frac{1}{N} \sum_{n \in [N]} \left[\mu \mathbf{y}_n^\top \tilde{\mathbf{z}}_n + (1 - \mu) \mathbf{y}_n^\top \mathbf{f}(\tilde{\mathbf{x}}_n; \mathbf{w}) \right] \right\} \\ &= \min_{\mathbf{w}} \max_{\mathbf{y}_n \in \mathcal{Y}, n \in [N], \alpha \geq 0} \left\{ \alpha \left(L(\mathbf{w}) - \sigma - \rho \right) - \frac{1}{N} \sum_{n \in [N]} \left[\mu \mathbf{y}_n^\top \tilde{\mathbf{z}}_n + (1 - \mu) \mathbf{y}_n^\top \mathbf{f}(\tilde{\mathbf{x}}_n; \mathbf{w}) \right] \right\} \\ &= \min_{\mathbf{w}} \frac{1}{N} \sum_{n \in [N]} \max_{\mathbf{y}_n \in \mathcal{Y}, \alpha \geq 0} \left\{ \alpha \ell(\tilde{\mathbf{z}}_n; \mathbf{f}(\tilde{\mathbf{x}}_n; \mathbf{w})) - \mu \mathbf{y}_n^\top \tilde{\mathbf{z}}_n - (1 - \mu) \mathbf{y}_n^\top \mathbf{f}(\tilde{\mathbf{x}}_n; \mathbf{w}) \right\} - \alpha(\sigma + \rho). \end{aligned} \quad (\text{EC1.3})$$

Notice that $\mathbf{w} = \tilde{\mathbf{w}}$ is always a feasible solution by assumption. Moreover, $\rho > 0$, hence $\mathbf{w} = \tilde{\mathbf{w}}$ is an interior point in the feasibility region. Additionally, the objective function of the inner maximization problem permits convexity in $\mathbf{w}$ and concavity in $(\mathbf{y}_n, \alpha)$. This achieves Slater's condition, and therefore strong duality holds. It is easy to see that the solution of Problem (EC1.2) coincides with that of Problem (EC1.3) when letting $\alpha^* = 1/\lambda$.

REMARK 1. Here, we make a quick comment about the limiting case $\mu = 1$. Notice that if $\mu = 1$, the regularization term does not involve the weights $\mathbf{w}$ in any form, hence the argmin obtained would coincide with $\arg \min_{\mathbf{w}} L(\mathbf{w})$, which is just $\hat{\mathbf{w}}$. However, it is not clear, that given a fixed λ bounded away from 0, that as $\mu \nearrow 1$, it is necessary for $\mathbf{w}(\mu) \rightarrow \hat{\mathbf{w}}$ uniformly. This is because the behaviour of μ is asymptotic at $\mu = 1$; beyond $\mu = 1$, there are no consistent ways for ensuring that Lemma 1 holds. Indeed, we see from our numerical simulations that as μ gets closer to 1, we do not recover $\hat{\mathbf{w}}$, as long as λ does not also uniformly decrease to 0. $\square$

EC1.7. Proof of Proposition 4 and Corollary 1

We prove both the proposition and the corollary simultaneously.

DEFINITION 1 (ROBUSTNESS DDR). The Robustness Decision-driven Regularization (RDDR, for short) model is defined as the following problem for $\mu \in [0, 1]$:

$$\max \quad \rho$$

$$\begin{aligned} \text{s.t. } & \frac{1}{N} \sum_{n \in [N]} \left[\mu \mathbf{y}_n^\top \tilde{\mathbf{z}}_n + (1 - \mu) \mathbf{y}_n^\top \mathbf{f}(\tilde{\mathbf{x}}_n; \mathbf{w}) \right] \leq \tau & \forall \mathbf{w} \in \mathcal{U}(\rho) & \quad (\mathbf{RDDR}) \\ & \rho > 0; \mathbf{y}_n \in \mathcal{Y} & \forall n \in [N], \end{aligned}$$

Note that when $\mu = 0$, it recovers the **JERO-Like** model. Once again, denote $\sigma = L(\tilde{\mathbf{w}})$. Now for the robust counterpart of the first constraint, it can be written as:

$$\begin{aligned} \max_{\mathbf{w}} \quad & \frac{1}{N} \sum_{n \in [N]} \left[\mu \mathbf{y}_n^\top \tilde{\mathbf{z}}_n + (1 - \mu) \mathbf{y}_n^\top \mathbf{f}(\tilde{\mathbf{x}}_n; \mathbf{w}) \right] \\ \text{s.t. } \quad & L(\mathbf{w}) \leq \sigma + \rho, \end{aligned}$$

the dual of which is

$$\min_{\alpha \geq 0} \max_{\mathbf{w}} \quad \frac{1}{N} \sum_{n \in [N]} \left[\mu \mathbf{y}_n^\top \tilde{\mathbf{z}}_n + (1 - \mu) \mathbf{y}_n^\top \mathbf{f}(\tilde{\mathbf{x}}_n; \mathbf{w}) \right] + \alpha (\sigma + \rho - L(\mathbf{w})).$$

Note that $\mathbf{w} = \tilde{\mathbf{w}}$ is feasible by definition. Moreover, since $\rho > 0$, $\mathbf{w} = \tilde{\mathbf{w}}$ is an interior point. Hence, Slater's condition is achieved, and strong duality holds. Further, $\rho > 0$ suffices to consider only $\alpha > 0$. Hence, using $\alpha = 1/\beta$ where $\beta > 0$,

$$\begin{aligned} \min_{\alpha > 0} \max_{\mathbf{w}} \quad & \frac{1}{N} \sum_{n \in [N]} \left[\mu \mathbf{y}_n^\top \tilde{\mathbf{z}}_n + (1 - \mu) \mathbf{y}_n^\top \mathbf{f}(\tilde{\mathbf{x}}_n; \mathbf{w}) \right] + \alpha (\sigma + \rho - L(\mathbf{w})) \leq \tau \\ \Leftrightarrow \min_{\frac{1}{\beta}} \max_{\mathbf{w}} \quad & \beta \frac{1}{N} \sum_{n \in [N]} \left[\mu \mathbf{y}_n^\top \tilde{\mathbf{z}}_n + (1 - \mu) \mathbf{y}_n^\top \mathbf{f}(\tilde{\mathbf{x}}_n; \mathbf{w}) \right] + \sigma + \rho - L(\mathbf{w}) - \beta \tau \leq 0 \end{aligned}$$

it follows that

$$\rho \leq \max_{\frac{1}{\beta}} \min_{\mathbf{w}} -\beta \frac{1}{N} \sum_{n \in [N]} \left[\mu \mathbf{y}_n^\top \tilde{\mathbf{z}}_n + (1 - \mu) \mathbf{y}_n^\top \mathbf{f}(\tilde{\mathbf{x}}_n; \mathbf{w}) \right] - \sigma + L(\mathbf{w}) + \beta \tau.$$

This allows us to reformulate the **1** problem as:

$$\max_{\mathbf{y}_n \in \mathcal{Y}, n \in [N], \frac{1}{\beta}} \min_{\mathbf{w}} -\beta \frac{1}{N} \sum_{n \in [N]} \left[\mu \mathbf{y}_n^\top \tilde{\mathbf{z}}_n + (1 - \mu) \mathbf{y}_n^\top \mathbf{f}(\tilde{\mathbf{x}}_n; \mathbf{w}) \right] - \sigma + L(\mathbf{w}) + \beta \tau.$$

On the other hand, by Theorem 1, **DDR** can be reformulated as:

$$\max_{\mathbf{y}_n \in \mathcal{Y}, n \in [N]} \min_{\mathbf{w}} \left\{ L(\mathbf{w}) + \lambda \frac{1}{N} \sum_{n \in [N]} \left[\tau - \mu \mathbf{y}_n^\top \tilde{\mathbf{z}}_n + (1 - \mu) \mathbf{y}_n^\top \mathbf{f}(\tilde{\mathbf{x}}_n; \mathbf{w}) \right] \right\}.$$

It follows that by setting $\lambda = \beta^*$ —the optimal β attained under the reformulated **RDDR**—the solution of **RDDR** coincides with that of **DDR**. $\square$

REMARK 2. Theorem 1 guarantees that, out-of-sample, if the optimal objective value of Robust-DDR is Γ^* , then $v_\mu(\tilde{\mathbf{w}}) \leq \Gamma^*$, with probability $\mathbb{P}[\tilde{\mathbf{w}} \in \mathcal{U}(\rho)]$, is related to the power of the learning model if L is the log-likelihood function. This is almost as good as saying that the out-of-sample costs are upper bounded with high probability, until we recall that the valuation function is only a surrogate for the true cost function.

EC1.8. Proof of Theorem 2

Define

$$\mathcal{Z}^*(\mathbf{w}) = \left\{ \tilde{\mathbf{z}} := \{\tilde{\mathbf{z}}_n \in \mathbb{R}^s\}_{n \in [N]} \left| \begin{array}{l} \exists \mathbf{y}_n \neq \mathbf{y}^*(\tilde{\mathbf{z}}_n), n \in [N]: \\ \left| \frac{1}{N} \sum_{n \in [N]} [\ell(\tilde{\mathbf{z}}_n; \tilde{\mathbf{z}}_n) - \ell(\tilde{\mathbf{z}}_n; \mathbf{f}(\tilde{\mathbf{x}}_n; \mathbf{w}))] \right| \leq t \\ |\mathbf{y}_n^\top \tilde{\mathbf{z}}_n - \mathbf{y}_n^\top \tilde{\mathbf{z}}_n| \leq \phi \quad \forall n \in [N] \\ |\mathbf{y}_n^\top \tilde{\mathbf{z}}_n - \mathbf{y}_n^\top \mathbf{f}(\tilde{\mathbf{x}}_n; \mathbf{w})| \leq \psi \quad \forall n \in [N] \\ |\mathbf{y}^*(\tilde{\mathbf{z}}_n)^\top \tilde{\mathbf{z}}_n - \mathbf{y}^*(\tilde{\mathbf{z}}_n)^\top \tilde{\mathbf{z}}_n| \leq \phi \quad \forall n \in [N] \\ |\mathbf{y}^*(\tilde{\mathbf{z}}_n)^\top \tilde{\mathbf{z}}_n - \mathbf{y}^*(\tilde{\mathbf{z}}_n)^\top \mathbf{f}(\tilde{\mathbf{x}}_n; \mathbf{w})| \leq \psi \quad \forall n \in [N] \\ |\mathbf{y}^*(\tilde{\mathbf{z}}_n)^\top \tilde{\mathbf{z}}_n - \mathbf{y}^*(\tilde{\mathbf{z}}_n)^\top \tilde{\mathbf{z}}_n| \leq \eta \quad \forall n \in [N] \end{array} \right. \right\}.$$

Notice that $\mathcal{Z}^*(\mathbf{w})$ is just $\mathcal{Z}(\mathbf{w})$, where we preserve the global condition and then demand that good estimates of the cost function are satisfied on two decision points, $\mathbf{y}^*(\mathbf{z}_n)$ and $\mathbf{y}_n$, which will be the optimal recourse solution obtained for the optimization problem at the end. As such, $\mathcal{Z}^*(\mathbf{w})$ relaxes the conditions for $\mathbf{z}_n$, thus $\mathcal{Z}^*(\mathbf{w}) \supseteq \mathcal{Z}(\mathbf{w})$, and the supremum over the latter is bounded above by the supremum of the former. Henceforth, we are interested in bounding the supremum over $\mathcal{Z}^*(\mathbf{w})$.

For ease of notations, hereinafter in this proof, we let $\hat{\mathbf{y}}_n \in \arg \min_{\mathbf{y} \in \mathcal{Y}} \mathbf{y}^\top \mathbf{f}(\tilde{\mathbf{x}}_n; \mathbf{w})$, $\check{\mathbf{y}}_n \in \arg \min_{\mathbf{y} \in \mathcal{Y}} \mathbf{y}^\top \tilde{\mathbf{z}}_n$, and $\tilde{\mathbf{y}}_n \in \arg \min_{\mathbf{y} \in \mathcal{Y}} \mathbf{y}^\top \tilde{\mathbf{z}}_n$ for $n \in [N]$. In what ensues, we shall bound the supremum by decomposing $\mathcal{Z}^*(\mathbf{w})$ into two steps. Denote

$$\mathcal{Z}_0(\mathbf{w}) = \left\{ \tilde{\mathbf{z}} := \{\tilde{\mathbf{z}}_n \in \mathbb{R}^s\}_{n \in [N]} \left| \begin{array}{l} |\check{\mathbf{y}}_n^\top \tilde{\mathbf{z}}_n - \tilde{\mathbf{y}}_n^\top \tilde{\mathbf{z}}_n| \leq \phi \quad \forall n \in [N] \\ |\tilde{\mathbf{y}}_n^\top \tilde{\mathbf{z}}_n - \tilde{\mathbf{y}}_n^\top \mathbf{f}(\tilde{\mathbf{x}}_n; \mathbf{w})| \leq \psi \quad \forall n \in [N] \\ |\tilde{\mathbf{y}}_n^\top \tilde{\mathbf{z}}_n - \tilde{\mathbf{y}}_n^\top \tilde{\mathbf{z}}_n| \leq \eta \quad \forall n \in [N] \end{array} \right. \right\},$$

and

$$\mathcal{Z}_1(\mathbf{w}) = \left\{ \tilde{\mathbf{z}} := \{\tilde{\mathbf{z}}_n \in \mathbb{R}^s\}_{n \in [N]} \left| \begin{array}{l} \exists \mathbf{y}_n \neq \check{\mathbf{y}}_n, n \in [N]: \\ |\mathbf{y}_n^\top \tilde{\mathbf{z}}_n - \mathbf{y}_n^\top \tilde{\mathbf{z}}_n| \leq \phi \quad \forall n \in [N] \\ |\mathbf{y}_n^\top \tilde{\mathbf{z}}_n - \mathbf{y}_n^\top \mathbf{f}(\tilde{\mathbf{x}}_n; \mathbf{w})| \leq \psi \quad \forall n \in [N] \\ \left| \frac{1}{N} \sum_{n \in [N]} [\ell(\tilde{\mathbf{z}}_n; \tilde{\mathbf{z}}_n) - \ell(\tilde{\mathbf{z}}_n; \mathbf{f}(\tilde{\mathbf{x}}_n; \mathbf{w}))] \right| \leq t \end{array} \right. \right\},$$

such that $\mathcal{Z}^*(\mathbf{w}) = \mathcal{Z}_0(\mathbf{w}) \cap \mathcal{Z}_1(\mathbf{w})$. In other words, the worst case regret is:

$$\begin{aligned} & \max_{\tilde{\mathbf{z}}_n \in \mathcal{Z}(\mathbf{w})} \left\{ \frac{1}{N} \sum_{n \in [N]} [\hat{\mathbf{y}}_n^\top \tilde{\mathbf{z}}_n - \check{\mathbf{y}}_n^\top \tilde{\mathbf{z}}_n] \right\} \\ & \leq \max_{\tilde{\mathbf{z}}_n \in \mathcal{Z}^*(\mathbf{w})} \left\{ \frac{1}{N} \sum_{n \in [N]} [\hat{\mathbf{y}}_n^\top \tilde{\mathbf{z}}_n - \check{\mathbf{y}}_n^\top \tilde{\mathbf{z}}_n] \right\} \\ & = \max_{\tilde{\mathbf{z}}_n \in \mathcal{Z}_0(\mathbf{w}) \cap \mathcal{Z}_1(\mathbf{w})} \left\{ \frac{1}{N} \sum_{n \in [N]} [\hat{\mathbf{y}}_n^\top \tilde{\mathbf{z}}_n - \check{\mathbf{y}}_n^\top \tilde{\mathbf{z}}_n] \right\} \\ & = \max_{\tilde{\mathbf{z}}_n \in \mathcal{Z}_1(\mathbf{w})} \left\{ \frac{1}{N} \sum_{n \in [N]} [\hat{\mathbf{y}}_n^\top \tilde{\mathbf{z}}_n - \check{\mathbf{y}}_n^\top \tilde{\mathbf{z}}_n] \right\} \end{aligned} \tag{EC1.4}$$

$$\begin{aligned} \text{s.t. } & \left| \tilde{\mathbf{y}}_n^\top \tilde{\mathbf{z}}_n - \tilde{\mathbf{y}}_n^\top \tilde{\mathbf{z}}_n \right| \leq \phi & \forall n \in [N]; \\ & \left| \tilde{\mathbf{y}}_n^\top \tilde{\mathbf{z}}_n - \tilde{\mathbf{y}}_n^\top \mathbf{f}(\tilde{\mathbf{x}}_n; \mathbf{w}) \right| \leq \psi & \forall n \in [N]; \\ & \left| \tilde{\mathbf{y}}_n^\top \tilde{\mathbf{z}}_n - \tilde{\mathbf{y}}_n^\top \tilde{\mathbf{z}}_n \right| \leq \eta & \forall n \in [N]. \end{aligned}$$

Here, we consider the Lagrangian relaxation of (EC1.4). Let $\boldsymbol{\alpha} = (\alpha^1, \alpha^2)$, $\boldsymbol{\beta} = (\beta^1, \beta^2)$, $\boldsymbol{\gamma} = (\gamma_n^1, \gamma_n^2)$ be the Lagrange multipliers. Then,

$$\begin{aligned} & \max_{\tilde{\mathbf{z}} \in \mathcal{Z}^*(\mathbf{w})} \left\{ \frac{1}{N} \sum_{n \in [N]} \left[\hat{\mathbf{y}}_n^\top \tilde{\mathbf{z}}_n - \tilde{\mathbf{y}}_n^\top \tilde{\mathbf{z}}_n \right] \right\} \\ & \leq \min_{\boldsymbol{\alpha}, \boldsymbol{\beta}, \boldsymbol{\gamma} \geq \mathbf{0}} \max_{\tilde{\mathbf{z}} \in \mathcal{Z}_1(\mathbf{w})} \left\{ \frac{1}{N} \sum_{n \in [N]} \left[\hat{\mathbf{y}}_n^\top \tilde{\mathbf{z}}_n - \tilde{\mathbf{y}}_n^\top \tilde{\mathbf{z}}_n \right. \right. \end{aligned} \quad (\text{EC1.5})$$

$$\begin{aligned} & \left. + \alpha_n^1 \left(\phi - \tilde{\mathbf{y}}_n^\top \tilde{\mathbf{z}}_n + \tilde{\mathbf{y}}_n^\top \tilde{\mathbf{z}}_n \right) + \alpha_n^2 \left(\phi + \tilde{\mathbf{y}}_n^\top \tilde{\mathbf{z}}_n - \tilde{\mathbf{y}}_n^\top \tilde{\mathbf{z}}_n \right) \right. \\ & \left. + \beta_n^1 \left(\psi - \tilde{\mathbf{y}}_n^\top \tilde{\mathbf{z}}_n + \tilde{\mathbf{y}}_n^\top \mathbf{f}(\tilde{\mathbf{x}}_n; \mathbf{w}) \right) + \beta_n^2 \left(\psi + \tilde{\mathbf{y}}_n^\top \tilde{\mathbf{z}}_n - \tilde{\mathbf{y}}_n^\top \mathbf{f}(\tilde{\mathbf{x}}_n; \mathbf{w}) \right) \right. \\ & \left. + \gamma_n^1 \left(\eta - \tilde{\mathbf{y}}_n^\top \tilde{\mathbf{z}}_n + \tilde{\mathbf{y}}_n^\top \tilde{\mathbf{z}}_n \right) + \gamma_n^2 \left(\eta + \tilde{\mathbf{y}}_n^\top \tilde{\mathbf{z}}_n - \tilde{\mathbf{y}}_n^\top \tilde{\mathbf{z}}_n \right) \right] \Big\} \\ & = \min_{\boldsymbol{\alpha}, \boldsymbol{\beta}, \boldsymbol{\gamma} \geq \mathbf{0}} \max_{\tilde{\mathbf{z}} \in \mathcal{Z}_1(\mathbf{w})} \left\{ \frac{1}{N} \sum_{n \in [N]} \left[\hat{\mathbf{y}}_n^\top \tilde{\mathbf{z}}_n + (\alpha_n^2 - \alpha_n^1 + \beta_n^2 - \beta_n^1 - 1) \tilde{\mathbf{y}}_n^\top \tilde{\mathbf{z}}_n \right. \right. \\ & \left. + (\alpha_n^1 - \alpha_n^2 - \gamma_n^1 + \gamma_n^2) \tilde{\mathbf{y}}_n^\top \tilde{\mathbf{z}}_n + (\beta_n^2 - \beta_n^1) \left(\hat{\mathbf{y}}_n^\top \mathbf{f}(\tilde{\mathbf{x}}_n; \mathbf{w}) - \tilde{\mathbf{y}}_n^\top \mathbf{f}(\tilde{\mathbf{x}}_n; \mathbf{w}) \right) \right. \\ & \left. - (\beta_n^2 - \beta_n^1) \hat{\mathbf{y}}_n^\top \mathbf{f}(\tilde{\mathbf{x}}_n; \mathbf{w}) + (\gamma_n^1 - \gamma_n^2) \tilde{\mathbf{y}}_n^\top \tilde{\mathbf{z}}_n + (\alpha_n^1 + \alpha_n^2) \phi + (\beta_n^1 + \beta_n^2) \psi + (\gamma_n^1 + \gamma_n^2) \eta \right] \Big\} \\ & \leq \max_{\tilde{\mathbf{z}} \in \mathcal{Z}_1(\mathbf{w})} \left\{ \frac{1}{N} \sum_{n \in [N]} \left[\hat{\mathbf{y}}_n^\top \tilde{\mathbf{z}}_n - 2 \hat{\mathbf{y}}_n^\top \mathbf{f}(\tilde{\mathbf{x}}_n; \mathbf{w}) \right] + a_0 \right\} \end{aligned} \quad (\text{EC1.6})$$

$$\leq \max_{\tilde{\mathbf{z}} \in \mathcal{Z}_1(\mathbf{w}); \mathbf{y}_n \in \mathcal{Y}, n \in [N]} \left\{ \frac{1}{N} \sum_{n \in [N]} \left[\mathbf{y}_n^\top \tilde{\mathbf{z}}_n - 2 \mathbf{y}_n^\top \mathbf{f}(\tilde{\mathbf{x}}_n; \mathbf{w}) \right] \right\} + a_0, \quad (\text{EC1.7})$$

where $a_0 = \phi + 2\psi + \eta + \frac{1}{N} \sum_{n \in [N]} \tilde{\mathbf{y}}_n^\top \tilde{\mathbf{z}}_n$. Here, inequality (EC1.5) follows from weak duality. Inequality (EC1.6) follows from two facts. First, we make the following choice of dual variables

$$(\alpha_n^1, \alpha_n^2, \beta_n^1, \beta_n^2, \gamma_n^1, \gamma_n^2) = (1, 0, 0, 2, 1, 0) \quad n \in [N].$$

Second, note that $\hat{\mathbf{y}}_n^\top \mathbf{f}(\tilde{\mathbf{x}}_n; \mathbf{w}) - \tilde{\mathbf{y}}_n^\top \mathbf{f}(\tilde{\mathbf{x}}_n; \mathbf{w}) \leq 0$, as $\hat{\mathbf{y}}_n \in \arg \min_{\mathbf{y}_n} \mathbf{y}_n^\top \mathbf{f}(\tilde{\mathbf{x}}_n; \mathbf{w})$ by definition. Finally, inequality (EC1.7) follows trivially from the fact that $\hat{\mathbf{y}}_n \in \mathcal{Y}$ for all $n \in [N]$.

Having relaxed $\hat{\mathbf{y}}_n$ to $\mathbf{y}_n$, we once again perform another Lagrangian relaxation, this time on $\mathcal{Z}_1(\mathbf{w})$. Let $\bar{\boldsymbol{\alpha}} = (\bar{\alpha}^1, \bar{\alpha}^2)$, $\bar{\boldsymbol{\beta}} = (\bar{\beta}^1, \bar{\beta}^2)$, $\bar{\boldsymbol{\gamma}} = (\bar{\gamma}^1, \bar{\gamma}^2)$ be the Lagrange multipliers. It follows that:

$$\begin{aligned} & \max_{\tilde{\mathbf{z}} \in \mathcal{Z}_1(\mathbf{w}); \mathbf{y}_n \in \mathcal{Y}, n \in [N]} \left\{ \frac{1}{N} \sum_{n \in [N]} \left[\mathbf{y}_n^\top \tilde{\mathbf{z}}_n - 2 \mathbf{y}_n^\top \mathbf{f}(\tilde{\mathbf{x}}_n; \mathbf{w}) \right] \right\} \\ & \leq \min_{\bar{\boldsymbol{\alpha}}, \bar{\boldsymbol{\beta}}, \bar{\boldsymbol{\gamma}} \geq \mathbf{0}} \max_{\tilde{\mathbf{z}} \in \mathcal{Z}_1(\mathbf{w}); \mathbf{y}_n \in \mathcal{Y}, n \in [N]} \left\{ \frac{1}{N} \sum_{n \in [N]} \left\{ \mathbf{y}_n^\top \tilde{\mathbf{z}}_n - 2 \mathbf{y}_n^\top \mathbf{f}(\tilde{\mathbf{x}}_n; \mathbf{w}) \right. \right. \\ & \left. + \bar{\alpha}_n^1 \left(\phi - \mathbf{y}_n^\top \tilde{\mathbf{z}}_n + \mathbf{y}_n^\top \tilde{\mathbf{z}}_n \right) + \bar{\alpha}_n^2 \left(\phi + \mathbf{y}_n^\top \tilde{\mathbf{z}}_n - \mathbf{y}_n^\top \tilde{\mathbf{z}}_n \right) \right. \\ & \left. + \bar{\beta}_n^1 \left(\psi - \mathbf{y}_n^\top \tilde{\mathbf{z}}_n + \mathbf{y}_n^\top \mathbf{f}(\tilde{\mathbf{x}}_n; \mathbf{w}) \right) + \bar{\beta}_n^2 \left(\psi + \mathbf{y}_n^\top \tilde{\mathbf{z}}_n - \mathbf{y}_n^\top \mathbf{f}(\tilde{\mathbf{x}}_n; \mathbf{w}) \right) \right\} \end{aligned} \quad (\text{EC1.8})$$

$$\begin{aligned}
& + \bar{\gamma}^1 \left(t - \frac{1}{N} \sum_{n \in [N]} [\ell(\tilde{\mathbf{z}}_n; \tilde{\mathbf{z}}_n) - \ell(\mathbf{f}(\tilde{\mathbf{x}}_n; \mathbf{w}); \tilde{\mathbf{z}})] \right) + \bar{\gamma}^2 \left(t + \frac{1}{N} \sum_{n \in [N]} [\ell(\tilde{\mathbf{z}}_n; \tilde{\mathbf{z}}_n) - \ell(\mathbf{f}(\tilde{\mathbf{x}}_n; \mathbf{w}); \tilde{\mathbf{z}})] \right) \Big\} \\
& = \min_{\bar{\alpha}, \bar{\beta}, \bar{\gamma} \geq 0} \max_{\mathbf{z}_n \in \mathcal{Z}; \mathbf{y}_n \in \mathcal{Y}, n \in [N]} \left\{ \frac{1}{N} \sum_{n \in [N]} \left[(1 - \bar{\alpha}_n^1 + \bar{\alpha}_n^2 - \bar{\beta}_n^1 + \bar{\beta}_n^2) \mathbf{y}_n^\top \tilde{\mathbf{z}}_n + (\bar{\alpha}_n^1 - \bar{\alpha}_n^2) \mathbf{y}_n^\top \tilde{\mathbf{z}}_n \right. \right. \\
& \quad \left. \left. + (\bar{\beta}_n^1 - \bar{\beta}_n^2 - 2) \mathbf{y}_n^\top \mathbf{f}(\tilde{\mathbf{x}}_n; \mathbf{w}) + (\bar{\gamma}^1 - \bar{\gamma}^2) (\ell(\mathbf{f}(\tilde{\mathbf{x}}_n; \mathbf{w}); \tilde{\mathbf{z}}_n) - \ell(\tilde{\mathbf{z}}_n; \tilde{\mathbf{z}}_n)) \right] \right\} \\
& \quad + \frac{1}{N} \sum_{n \in [N]} \left[(\bar{\alpha}_n^1 + \bar{\alpha}_n^2) \phi + (\bar{\beta}_n^1 + \bar{\beta}_n^2) \psi \right] + (\bar{\gamma}^1 + \bar{\gamma}^2) t \\
& \leq \bar{\lambda} \frac{1}{N} \sum_{n \in [N]} \ell(\mathbf{f}(\tilde{\mathbf{x}}_n; \mathbf{w}); \tilde{\mathbf{z}}_n) + \frac{1}{N} \sum_{n \in [N]} \max_{\mathbf{y}_n \in \mathcal{Y}} \left[-\mu \mathbf{y}_n^\top \tilde{\mathbf{z}}_n - (1 - \mu) \mathbf{y}_n^\top \mathbf{f}(\tilde{\mathbf{x}}_n; \mathbf{w}) \right] + a_1, \tag{EC1.9}
\end{aligned}$$

where $a_1 = \phi + 2\psi + t$. The inequality (EC1.8) is once again due to the weak duality, and inequality (EC1.9) is because of the triangle inequality assumption on $\ell(\mathbf{u}; \mathbf{v})$ and the following choice of dual variables,

$$(\bar{\alpha}_n^1, \bar{\alpha}_n^2, \bar{\beta}_n^1, \bar{\beta}_n^2, \bar{\gamma}^1, \bar{\gamma}^2) = \left(\frac{1-\mu}{2}, \frac{1+\mu}{2}, \frac{3+\mu}{2}, \frac{1-\mu}{2}, \frac{1+\bar{\lambda}}{2}, \frac{1-\bar{\lambda}}{2} \right) \quad n \in [N]$$

for some $\mu \in [-1, 1]$ and $\bar{\lambda} \in (0, 1]$.

Therefore, we have

$$\begin{aligned}
& \max_{\tilde{\mathbf{z}} \in \mathcal{Z}(\mathbf{w})} \left\{ \frac{1}{N} \sum_{n \in [N]} \left[\mathbf{y}^\star(\tilde{\mathbf{z}}_n)^\top \tilde{\mathbf{z}}_n - \min_{\mathbf{y}_n \in \mathcal{Y}} \tilde{\mathbf{y}}_n^\top \tilde{\mathbf{z}}_n \right] \right\} \\
& \leq \max_{\tilde{\mathbf{z}} \in \mathcal{Z}_1(\mathbf{w}); \mathbf{y}_n \in \mathcal{Y}, n \in [N]} \left\{ \frac{1}{N} \sum_{n \in [N]} \left[\mathbf{y}_n^\top \tilde{\mathbf{z}}_n - 2\mathbf{y}_n^\top \mathbf{f}(\tilde{\mathbf{x}}_n; \mathbf{w}) \right] + a_0 \right\} \\
& \leq \bar{\lambda} \frac{1}{N} \sum_{n \in [N]} \ell(\mathbf{f}(\tilde{\mathbf{x}}_n; \mathbf{w}); \tilde{\mathbf{z}}_n) - \frac{1}{N} \sum_{n \in [N]} \min_{\mathbf{y}_n \in \mathcal{Y}} \left[\mu \mathbf{y}_n^\top \tilde{\mathbf{z}}_n + (1 - \mu) \mathbf{y}_n^\top \mathbf{f}(\tilde{\mathbf{x}}_n; \mathbf{w}) \right] + a_0 + a_1 = \frac{1}{\lambda} V(\mathbf{w}) + a,
\end{aligned}$$

with $\lambda = 1/\bar{\lambda}$ and $a := a_0 + a_1 = 2\phi + 4\psi + \eta + t + \frac{1}{N} \sum_{n \in [N]} \tilde{\mathbf{y}}_n^\top \tilde{\mathbf{z}}_n$. Notice here that $\frac{1}{N} \sum_{n \in [N]} \tilde{\mathbf{y}}_n^\top \tilde{\mathbf{z}}_n$ is evaluated only over the training data $\tilde{\mathbf{z}}_n$ and hence is a constant given the training dataset.

Hence, in particular, if we choose $\mathbf{w} = \hat{\mathbf{w}}^{\text{DDR}}$, we have

$$\max_{\tilde{\mathbf{z}} \in \mathcal{Z}(\mathbf{w}^{\text{DDR}})} \left\{ \frac{1}{N} \sum_{n \in [N]} \left[\mathbf{y}^\star(\tilde{\mathbf{z}}_n)^\top \tilde{\mathbf{z}}_n - \min_{\mathbf{y}_n \in \mathcal{Y}} \mathbf{y}_n^\top \tilde{\mathbf{z}}_n \right] \right\} \leq \frac{1}{\lambda} V(\hat{\mathbf{w}}^{\text{DDR}}) + a.$$

This completes the proof. $\square$

EC1.9. Proof of Proposition 5

When $L(\mathbf{w}) = (1 - \mu) \frac{1}{N} \sum_{n \in [N]} \mathbf{y}^\star(\tilde{\mathbf{z}}_n)^\top \mathbf{f}(\tilde{\mathbf{x}}_n; \mathbf{w})$, $\lambda = 1$, and $\mu = -1$, **DDR** becomes

$$\min_{\mathbf{w}} (1 - \mu) \frac{1}{N} \sum_{n \in [N]} \mathbf{y}^\star(\tilde{\mathbf{z}}_n)^\top \mathbf{f}(\tilde{\mathbf{x}}_n; \mathbf{w}) - \frac{1}{N} \sum_{n \in [N]} \min_{\mathbf{y}_n \in \mathcal{Y}} \left\{ -\mathbf{y}_n^\top \tilde{\mathbf{z}}_n + 2\mathbf{y}_n^\top \mathbf{f}(\tilde{\mathbf{x}}_n; \mathbf{w}) \right\},$$

which is exactly the **SPO+** model. $\square$

REMARK 3. **SPO+** occurs when $\mu = -1$. In this case, **DDR** reduces to, for some $\lambda \geq 0$:

$$\min_{\mathbf{w}} L(\mathbf{w}) - \lambda \left(\frac{1}{N} \sum_{n \in [N]} \min_{\mathbf{y}_n \in \mathcal{Y}} \left\{ 2\mathbf{y}_n^\top \mathbf{f}(\tilde{\mathbf{x}}_n; \mathbf{w}) - \mathbf{y}_n^\top \tilde{\mathbf{z}}_n \right\} \right). \tag{EC1.10}$$

One way to understand the motivation of such a model, is to notice that the expression within the infimum has the following decomposition into the estimated cost and its estimation error when compared against the empirical cost:

$$\underbrace{\mathbf{y}^\top \mathbf{f}(\tilde{\mathbf{x}}_n, \mathbf{w}) - \mathbf{y}^\top \tilde{\mathbf{z}}}_{\text{estimation error}} + \underbrace{\mathbf{y}^\top \mathbf{f}(\tilde{\mathbf{x}}_n, \mathbf{w})}_{\text{best estimate}}. \quad (\text{EC1.11})$$

We remark here that if we let $\phi = 0, \eta = 0$ and t, ψ to be sufficiently large such that the first and the third conditions in $\mathcal{Z}(\mathbf{w})$ no longer exist, *i.e.*,

$$\mathcal{Z}(\mathbf{w}) = \left\{ \mathbf{z} := \{ \mathbf{z}_n \in \mathbb{R}^s \}_{n \in [N]} \left| \begin{array}{ll} \forall \mathbf{y}_n \in \mathcal{Y}, n \in [N] : & \\ \left| \mathbf{y}_n^\top \mathbf{z}_n - \mathbf{y}_n^\top \tilde{\mathbf{z}}_n \right| = 0 & \forall n \in [N] \\ \left| \mathbf{y}^*(\mathbf{z}_n)^\top \tilde{\mathbf{z}}_n - \mathbf{y}^*(\tilde{\mathbf{z}}_n)^\top \tilde{\mathbf{z}}_n \right| = 0 & \forall n \in [N] \end{array} \right. \right\},$$

Regret-DDR is exactly **SPO**. Furthermore, if one follows the proof for Theorem 2 with these specifications of the parameters, then one will also recover the construction of **SPO+** from **SPO**. This once again reinforces the notion that the construction of **DDR** as an approximation to **Regret-DDR** is as canonical as the construction of **SPO+** as an approximation to **SPO**.

EC2. Sub-optimality of the SLO approach and relationship of DDR to other regularizers

EC2.1. Sub-optimality of the SLO approach

To better understand why directly minimizing prediction error is sub-optimal, consider the toy knapsack problem of picking smallest cost $\mathbf{z}$, where cost function is $c(\mathbf{y}; \mathbf{z}) = \mathbf{y}^\top \mathbf{z}$, the feasible set is $\mathcal{Y} = \{ \mathbf{y} \in [0, 1]^N, \mathbf{y}^\top \mathbf{1} = 1 \}$, and $\mathbf{z}$ is generated from some true linear relationship $\mathbf{z} = \tilde{\mathbf{w}}^\top \mathbf{x} + \epsilon$, with true parameter $\tilde{\mathbf{w}}$. Typically, prediction error is measured globally. Thus, the residuals ϵ of a model that minimizes global prediction error, will be smaller over less extreme-valued regions of $\mathbf{x}$ and vice versa, as illustrated by the dark blue shaded confidence interval in Figure EC2.1 below. However, in the knapsack problem, we are concerned about extreme values of $\mathbf{z}$, which occur over the region of extreme values of $\mathbf{x}$ where the prediction error is currently largest. In other words, more generally, the cost minimization problem induces a region $\bar{\mathcal{X}} \subset \mathcal{X}$ over context $\mathbf{x}$, corresponding to a region $\bar{\mathcal{Z}} \subset \mathcal{Z}$ of the uncertainty $\mathbf{z}$, where optimal decisions to the cost minimization problem are more likely to occur than not. Moreover, this region $\bar{\mathcal{X}}$ is often extreme in relation to the support $\mathcal{X}$ as optimization, by definition, seeks extreme results. If one is able to improve the predictive accuracy over this region $\bar{\mathcal{Z}}$, then the optimization model is better able to differentiate between alternative decisions $\mathbf{y}$. Thus, a model with confidence interval indicated by the broken line in Figure EC2.1 is expected to have higher decision power than its shaded blue counterpart, as it incurs smaller prediction error over the region that matters—the small values of $\mathbf{z}$. It may have a larger overall prediction error, but that results from sacrificing prediction accuracy over non-critical regions of $\mathcal{X}$ —that lead to large values of $\mathbf{z}$ —in favour of $\bar{\mathcal{X}}$.

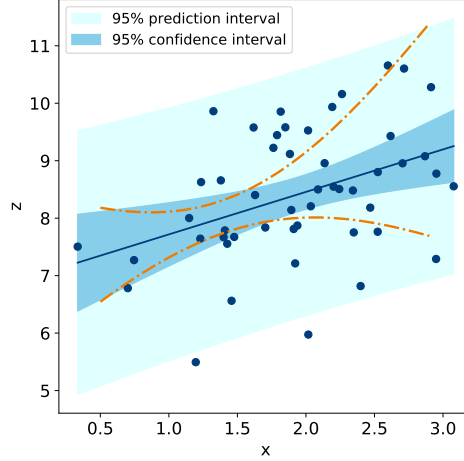


Figure EC2.1 One-dimensional illustrative of centred (dark blue shaded area) and non-centred (broken orange line) confidence intervals

EC2.2. Relationship to other regularizers

In the **DDR** model, the regularizer $R_\mu(\mathbf{w}) := -\frac{1}{N} \sum_n \nu_\mu^n(\mathbf{w})$ is chosen as the sample average approximation of the valuation function $\nu_\mu(\mathbf{w})$. Here, we shorthand $\nu_\mu^n(\mathbf{w})$ to mean the valuation function evaluated on the single data point indexed by n , $(\tilde{\mathbf{x}}, \tilde{\mathbf{z}}) = (\tilde{\mathbf{x}}_n, \tilde{\mathbf{z}}_n)$. In both Lasso and Ridge regressions, the regularizer is chosen as some norm: $R(\mathbf{w}) = \|\mathbf{w}\|_q$. Nonetheless, recall that each norm can be represented as the optimizer of some optimization problem. By the definition of dual norm, we have

$$R(\mathbf{w}) = \|\mathbf{w}\|_q = \max_{\boldsymbol{\xi}} \sum_{i \in [p]} \sum_{j \in [d]} w_{ij} \xi_{i,j}, \quad (\text{EC2.12})$$

$$\text{s.t. } \boldsymbol{\xi} \in \Xi := \{\boldsymbol{\xi} \in \mathbb{R}^{p \times d} : \|\boldsymbol{\xi}\|_{q^*} \leq 1\}$$

where $\|\cdot\|_{q^*}$ represents the dual norm of $\|\cdot\|_q$, and $\frac{1}{q} + \frac{1}{q^*} = 1$.

Consider the case where the learner is linear, *i.e.*, $\tilde{\mathbf{z}} := \mathbf{f}(\mathbf{x}; \mathbf{w}) = \mathbf{w}^\top \mathbf{x}$. In the case where $\mu = 0$, the decision-driven regularizer is

$$\begin{aligned} R_0(\mathbf{w}) &= \frac{1}{N} \sum_n \nu_0^n(\mathbf{w}) = \max_{\mathbf{y}_n \in \mathcal{Y}} -\frac{1}{N} \sum_n \mathbf{y}_n^\top \mathbf{w}^\top \mathbf{x}_n \\ &= \max_{\mathbf{y}_n \in \mathcal{Y}} \sum_{i \in [p]} \sum_{j \in [d]} w_{ij} \left(-\frac{1}{N} \sum_n x_{n,i} y_{n,j} \right). \end{aligned}$$

Thus, we may express $R_0(\mathbf{w})$ as the following:

$$\begin{aligned} R_0(\mathbf{w}) &= \max_{\boldsymbol{\xi}} \sum_{i \in [p]} \sum_{j \in [d]} w_{ij} \xi_{i,j}, \quad (\text{EC2.13}) \\ \text{s.t. } \boldsymbol{\xi} &\in \Xi_0^N := \left\{ \boldsymbol{\xi} \in \mathbb{R}^{p \times d} : \exists \mathbf{y}_n \in \mathcal{Y}, n \in [N], \xi_{i,j} = -\frac{1}{N} \sum_n x_{n,i} y_{n,j} \right\} \end{aligned}$$

that depends on the dataset $\mathcal{D}^N$. Denote X^N as the $p \times N$ matrix that collects all p -dimensional feature vectors in the dataset of N points, and denote $\mathcal{Y}^N := \mathcal{Y} \times \cdots \times \mathcal{Y}$ as the product space of $\mathcal{Y}$. Then, we can notice that $\Xi_0^N = \frac{1}{N} X^{N^\top} \mathcal{Y}^N := \{\frac{1}{N} X^{N^\top} \mathbf{y} : \mathbf{y} \in \mathcal{Y}^N\}$, which is the linear transformation of the feasibility space by the data matrix X^N .

The forms of (EC2.12) and (EC2.13) allow us to draw the connection between decision-driven regularization and norm-based regularizers like Lasso and Ridge regressions. We wish to analyze this relationship by assuming that we are currently at the OLS solution, *i.e.*, $\mathbf{w} = \hat{\mathbf{w}}^{\text{OLS}}$, and investigating what impact adding either of the regularizers has on the performance of the solution – specifically, how the base solution $\hat{\mathbf{w}}^{\text{OLS}}$, results in different weights, such as $\hat{\mathbf{w}}^{\text{Lasso}}$ in the case of Lasso. Here, we examine locally around the neighbourhood of OLS weights, $\hat{\mathbf{w}}^{\text{OLS}}$, the directions of improving (decreasing) $L(\mathbf{w}) + \lambda R(\mathbf{w})$ for different regularizers $R(\cdot)$. As $\hat{\mathbf{w}}^{\text{OLS}}$ is the minimizer of $L(\mathbf{w})$ (hence gradient of L at $\mathbf{w} = \hat{\mathbf{w}}^{\text{OLS}}$ is 0), the action of adding $R(\mathbf{w})$ is to draw the solution in the direction of the gradient of $R(\mathbf{w})$. Closer examination of problems (EC2.12) and (EC2.13) reveals that these directions are approximated by the optimal solutions of ξ at $\mathbf{w} = \hat{\mathbf{w}}^{\text{OLS}}$. This is because ξ are the dual variables in the convex conjugate forms of $R(\mathbf{w})$. In other words, around the neighborhood of base solution $\hat{\mathbf{w}}^{\text{OLS}}$, the base solution is pushed along the optimal ξ 's to problems (EC2.12) and (EC2.13) at $\mathbf{w} = \hat{\mathbf{w}}^{\text{OLS}}$ towards the norm-based regularizer and DDR weights, respectively. However, problems (EC2.12) and (EC2.13) have the same objective, and hence the same optimizing direction, but with different feasibility sets Ξ and Ξ_0^N . Thus, the eventual optimal solution of ξ depends on how different these feasibility sets are.

Our goal is to examine the impact on the performance of the eventual decisions, of pulling the OLS weights in the directions of the solutions to (EC2.12) and (EC2.13). If the solution ξ to problem (EC2.13) is a local improving direction for the performance, *i.e.*, that DDR improves on the performance of OLS on the eventual cost function, and if the angle between the two solutions ξ to problem (EC2.12) and (EC2.13) is fairly large, then we can conclude that the direction in which norm-based regularizers pull the OLS weights towards, *i.e.*, the solution to problem (EC2.12), may very likely be a local worsening direction for the performance of the eventual cost function.

This actually turns out to be the case in our simulations. We see that DDR improves over OLS, and that the angles of the local directions from DDR to Lasso and Ridge regressions are around 60°. In Table EC2.1, we present the summary statistics of these angles amongst 100 simulations. These angles are significantly large—in other words, the regularizers in Lasso and Ridge regressions are pulling the OLS weights away in very different directions from the decision-driven regularizer. As a consequence, we discover, ironically, that while Lasso and Ridge regressions both improve the prediction accuracy over OLS, they both lead to significantly poorer performance on the cost function.

Table EC2.1 Angles of local directions of Lasso and Ridge to DDR in 100 simulations

Angles to DDR	Lasso	Ridge
Mean	82°	72°
10th Percentile	80°	72°
90th Percentile	85°	73°

In conclusion, what we have explained here is that blindly adopting regularizers, even if for the purposes of improving prediction accuracy, might not lead to better performance. In fact, we can potentially

see the opposite, where they lead to poorer performance, pulling the optimal solution away from the direction of improving performance.

EC3. Pseudo-code for Solving RDDR

COROLLARY 1. *Under the same assumptions as in Theorem 1, then for all $\mu \in [0, 1)$ and τ for which $\exists \mathbf{y}_n \in \mathcal{Y}, n \in [N]$ such that the target is attained $\frac{1}{N} \sum_{n \in [N]} [\mu \mathbf{y}_n^\top \tilde{\mathbf{z}}_n + (1 - \mu)c(\mathbf{y}_n; \mathbf{f}(\tilde{\mathbf{x}}_n, \tilde{\mathbf{w}}))] \leq \tau$, under learning optimal weights $\tilde{\mathbf{w}} = \arg \min L(\mathbf{w})$, there exists some $\lambda := \lambda(\tau) \geq 0$ for which the solutions of RDDR and the worst-case weights attained under optimal uncertainty set size ρ^* from RDDR coincide.*

In short, we seek the best radius of the uncertainty set ρ , via bisection search, wherein we solve a feasibility sub-problem in each iteration. Here, we present the pseudo-code for the algorithm, however, for more details, the reader can be directed to [Zhu et al. \(2022\)](#), as the procedure for solving the model follows identically.

Algorithm 1 RDDR

Input τ and $\tilde{\mathbf{z}}$.
Initialization: $\rho_1 = 0, \rho_2 = \bar{\rho}$, where $\bar{\rho}$ is a sufficiently large number.
while $\rho_2 - \rho_1 > \epsilon$ **do**
 $\rho := (\rho_1 + \rho_2)/2$
 Solve the subproblem **Sub-RDDR**, obtain optimal value δ^* and optimal decisions $\mathbf{y}_n^*$.
 if $\delta^* > 0$ **then**
 $\rho_2 = \rho$
 else
 $\rho_1 = \rho$
 end if
end while
Output: Optimal $\rho^* = (\rho_1 + \rho_2)/2$ and optimal decision $\mathbf{y}_n^{\text{RDDR}} = \mathbf{y}_n^*$.
Input Optimal $\rho = \rho^*$ and optimal decision $\mathbf{y} = \mathbf{y}_n^*$
Do Solve the Problem (**ROBUSTW**), and obtain $\mathbf{w}^*$
Output Worst scenario $\tilde{\mathbf{w}}^{\text{RDDR}} = \mathbf{w}^*$

In this algorithm, the sub-problem is described as:

$$\begin{aligned} \min_t \quad & t - \tau & (\text{Sub-RDDR}) \\ \text{s.t.} \quad & \frac{1}{N} \sum_{n \in [N]} [\mu \mathbf{y}_n^\top \tilde{\mathbf{z}}_n + (1 - \mu) \mathbf{y}_n^\top \mathbf{f}(\tilde{\mathbf{x}}_n; \mathbf{w})] \leq t, & \forall \mathbf{w} \in \mathcal{U}(\rho) \\ \text{where} \quad & \mathcal{U}(\rho) := \{\mathbf{w} : L(\mathbf{w}) - L(\tilde{\mathbf{w}}) \leq \rho\}. \end{aligned}$$

We finally utilize ρ^* and $\mathbf{y}^*$ derived from the overarching problem to solve for the worst case $\tilde{\mathbf{w}}^{\text{RDDR}}$. This is done via,

$$\arg \max_{\mathbf{w} \in \mathcal{U}(\rho^*)} \frac{1}{N} \sum_{n \in [N]} [\mu c(\mathbf{y}_n^*, \tilde{\mathbf{z}}_n) + (1 - \mu) c(\mathbf{y}_n^*, \mathbf{f}(\tilde{\mathbf{x}}_n; \mathbf{w}))] \quad (\text{ROBUSTW})$$

EC4. Additional Simulation Results

In this Appendix, we present further results from the numerical illustration that were omitted from the main text for brevity. This includes the robustness of our findings to different parameters and a short study into the case when the network size is varied.

EC4.1. A Closer Look at Over-fitting in the Tree-based SLO Benchmarks

In §4.1 of the main text, Table 2 reports the in- and out-of-sample prediction accuracy of the SLO benchmarks (OLS, RF, XGBoost) and DDR under their default settings, which points to potential over-fitting in the two tree-based learners. To further elucidate the extent to which such over-fitting depends on model complexity, Table EC4.2 refines the entries for RF and XGBoost by sweeping the maximum tree depth across $\{2, 4, 6, 8, \text{Default}\}$, while keeping OLS and DDR unchanged as references. We can see that changing the hyperparameters might not resolve over-fitting, but is rather an inherent feature of the learner.

Table EC4.2 Prediction accuracy of SLO models in terms of root mean square error (RMSE) across multiple datasets

Models	In-sample RMSE			Out-of-sample RMSE		
	5 th %-tile	Mean	95 th %-tile	5 th %-tile	Mean	95 th %-tile
OLS	3.8393	3.9416	4.0682	4.1279	4.1846	4.2395
DDR	3.8418	3.9440	4.0705	4.1286	4.1830	4.2352
RF(Maximum depth = 2)	4.0068	4.1499	4.2974	4.3693	4.4779	4.5875
RF(Maximum depth = 4)	3.36	3.4682	3.5785	4.3047	4.3904	4.4816
RF(Maximum depth = 6)	2.4572	2.6123	2.782	4.3191	4.3939	4.4766
RF(Maximum depth = 8)	1.8514	1.9727	2.1223	4.338	4.4143	4.504
RF(Default)	1.5872	1.6531	1.7039	4.3479	4.4255	4.5155
XGBoost(Maximum depth = 2)	3.924	4.0634	4.221	4.3788	4.4943	4.6057
XGBoost(Maximum depth = 4)	3.38	3.5202	3.6778	4.4097	4.517	4.6255
XGBoost(Maximum depth = 6)	2.9835	3.114	3.2322	4.4822	4.5833	4.6913
XGBoost(Maximum depth = 8)	2.8146	2.9393	3.0729	4.5155	4.6173	4.7142
XGBoost(Default)	2.9835	3.1140	3.2322	4.4822	4.5833	4.6913

EC4.2. Robustness of Our Model Across Various Settings

This section presents a series of numerical experiments designed to evaluate the robustness of the proposed DDR model under varying conditions. Specifically, we investigate the influence of training data size, feature dimensionality, noise level, and model misspecification. To assess the impact of training data size, the sample size N is varied across the set $\{50, 100, 200, 500, 1000\}$, while all other parameters are held constant at their baseline values. Likewise, the number of features p is selected from $\{1, 3, 5, 7, 10, 15\}$. To evaluate the effect of model misspecification, the parameter β is varied over the set $\{0.4, 0.8, 1.0, 1.4, 1.6, 2.0, 3.0, 4.0\}$. Finally, the noise level $\bar{\epsilon}$ is examined by selecting values from $\{0.25, 0.5, 0.75, 1.0\}$. The comparison between DDR and OLS is summarized in Figure EC4.2.

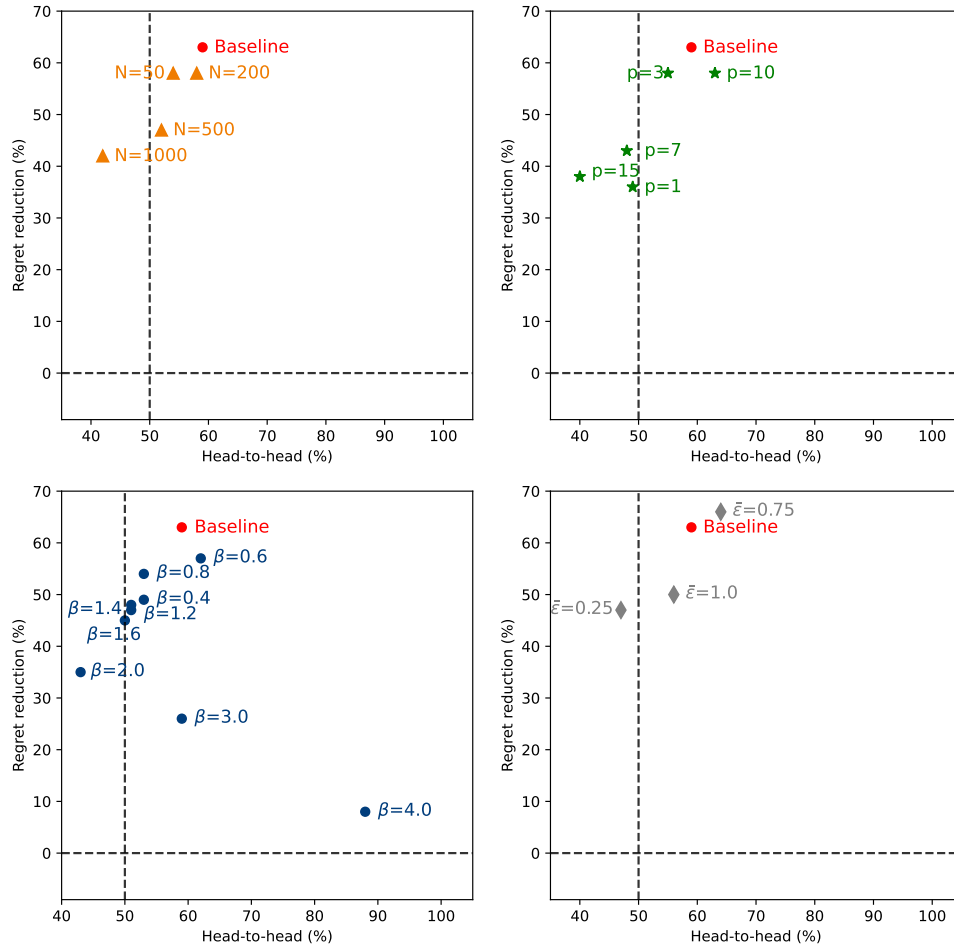


Figure EC4.2 Performance of DDR against OLS under 3×3 grid and $\mu = 0.75$ and $\lambda = 0.8$ under various settings

From Figure EC4.2, we first observe that the DDR model consistently attains superior regret performance, as indicated by positive regret reduction across all evaluated settings. This finding highlights the advantage of employing the DDR model. However, it is also apparent that in certain scenarios, the head-to-head performance of the DDR model does not surpass that of OLS. For instance, when assessing the impact of training data size, the DDR model exhibits improved head-to-head performance across all sample sizes except at $N = 1000$. This outcome is anticipated, as the OLS estimator's accuracy improves with larger sample sizes, thereby reducing the performance gap between DDR and OLS in such cases. Nevertheless, Figure EC4.2 demonstrates that our DDR model outperforms in terms of both regret reduction and head-to-head metrics in the majority of cases.

Similarly, the comparison between the DDR and SPO+ approaches is presented in Figure EC4.3. The results indicate that the DDR model consistently outperforms SPO+ across all settings in terms of both regret reduction and head-to-head performance metrics, thereby demonstrating the superiority of the DDR model over the SPO+ baseline.

EC5. DDR Tree

While we have elected to use a linear learner in our numerical results, DDR can more generally be applied to other learners. To illustrate this, in this section, we consider a tree-based learner in DDR, and term it as DDR tree.

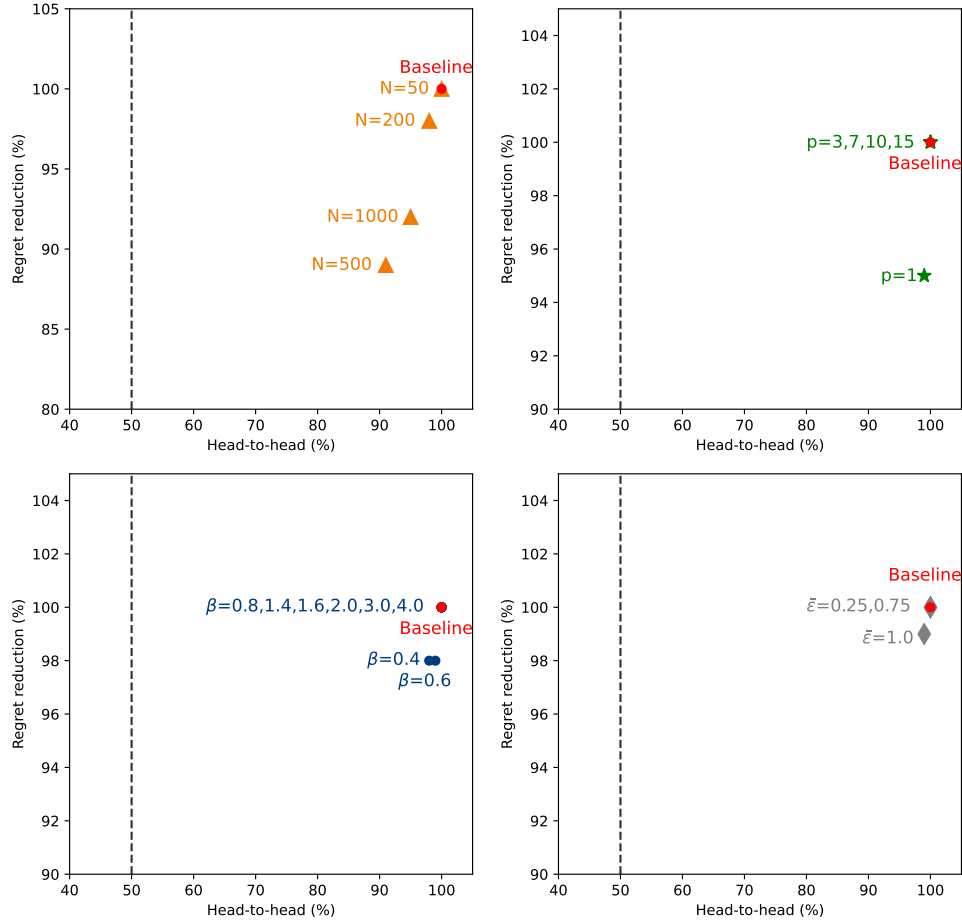


Figure EC4.3 Performance of DDR against SPO+ under 3×3 grid and $\mu = 0.75$ and $\lambda = 0.8$ under various settings

Consider a decision tree with K leaf nodes, where the predictions $\hat{z}$ are constrained to be equal if they fall into the same leaf node. We model this by defining R_k the set of samples assigned to leaf node k , and predictions $\hat{z}_k$ as being indexed by the leaf nodes $k \in [K]$. We write the DDR loss objective function as:

$$L_{\text{DDR}}(R_k, \hat{z}_k) = \sum_{k \in [K]} \sum_{n \in R_k} \|\tilde{z}_n - \hat{z}_k\|_2^2 + \frac{\lambda}{N} \sum_{k \in [K]} \sum_{n \in R_k} \max_{\mathbf{y}_n \in \mathcal{Y}} \{ -\mu \mathbf{y}_n^\top \tilde{z}_n - (1 - \mu) \mathbf{y}_n^\top \hat{z}_k \}$$

where the first term represents the prediction accuracy and the second term represents the decision cost. Here, sets R_k and predictions $\hat{z}_k$ are the parameters (“weights”) of the tree to be learned and determined.

Using Proposition 3, the DDR loss function can be further reformulated as:

$$L_{\text{DDR}}(R_k, \hat{z}_k) = \frac{\lambda}{N} \sum_{k \in [K]} \sum_{n \in R_k} \left(\frac{N}{\lambda} \|\tilde{z}_n - \hat{z}_k\|_2^2 + \min_{\mathbf{A}^\top \boldsymbol{\kappa}_n \geq -\mu \tilde{z}_n - (1 - \mu) \hat{z}_k} \boldsymbol{\kappa}_n^\top \mathbf{b} \right).$$

Accordingly, the prediction cost vectors for the leaf nodes, given the partitioning of observations into leaf nodes, *i.e.*, the sets R_k , are determined by solving the following constrained optimization problem:

$$\hat{\mathbf{z}}^{\text{DDR}} = \arg \min_{\hat{\mathbf{z}}_k : \forall k \in [K]} \left\{ \frac{\lambda}{N} \sum_{k \in [K]} \sum_{n \in R_k} \left(\frac{N}{\lambda} \|\tilde{z}_n - \hat{z}_k\|_2^2 + \boldsymbol{\kappa}_n^\top \mathbf{b} \right) \middle| \mathbf{A}^\top \boldsymbol{\kappa}_n \geq -\mu \tilde{z}_n - (1 - \mu) \hat{z}_k, \forall n \in R_k, k \in [K] \right\}$$

If the sets R_k are known, the above optimization problem is easy. Instead, we propose the following two methodologies to optimize for R_k .

EC5.1. Recursive Partitioning Approach.

Emulating the CART algorithm, we construct the decision tree via a greedy recursive partitioning strategy. Let x_{nj} denote the j^{th} feature component of the n^{th} observation in the training data. Starting from the entire training dataset, we consider binary splits defined by pairs (j, s) , where j indicates the selected feature and s denotes the split threshold. This split partitions the data into two disjoint subsets:

$$R_1(j, s) = \{n \in [N] \mid x_{nj} \leq s\} \quad \text{and} \quad R_2(j, s) = \{n \in [N] \mid x_{nj} > s\}.$$

The optimal first split of the decision tree is determined by selecting the pair (j, s) that minimizes the total L_{DDR} loss over both subsets:

$$\begin{aligned} \min_{\hat{\mathbf{z}}_1, \hat{\mathbf{z}}_2} \quad & \frac{\lambda}{N} \left[\sum_{n \in R_1} \left(\frac{N}{\lambda} \|\tilde{\mathbf{z}}_n - \hat{\mathbf{z}}_1\|_2^2 + \boldsymbol{\kappa}_n^\top \mathbf{b} \right) + \sum_{n \in R_2} \left(\frac{N}{\lambda} \|\tilde{\mathbf{z}}_n - \hat{\mathbf{z}}_2\|_2^2 + \boldsymbol{\kappa}_n^\top \mathbf{b} \right) \right] \\ \text{s.t.} \quad & \mathbf{A}^\top \boldsymbol{\kappa}_n \geq -\mu \tilde{\mathbf{z}}_n - (1 - \mu) \hat{\mathbf{z}}_k, \quad \forall n \in R_k, k = 1, 2 \end{aligned}$$

Once the optimal split is selected, the procedure is then recursively applied greedily to obtain binary splits in the resulting leaves until the stopping criteria is met. The prediction cost vectors $\hat{\mathbf{z}}^{\text{DDR}}$ for each leaf node are then optimized simultaneously, after the complete tree structure has been determined.

EC5.2. Integer Programming Approach.

Motivated by recent advances in decision tree construction through optimization techniques (?), DDR tree can be solved as an integer programming model. Specifically, we aim to construct a decision tree based on a pre-specified tree depth H and a minimum number of training observations $N_{\min}$ required in each leaf node.

Note that not all leaf nodes are required to be active (*i.e.*, contain training observations), and similarly, not all branch nodes need to be active splits (*i.e.*, nodes that partition the data). Denote the index sets of leaf and branch nodes by $\mathcal{T}_K = \{1, \dots, K\}$ and $\mathcal{T}_B = \{1, \dots, B\}$, respectively. Let $g_k = \mathbb{I}\{\text{leaf } k \text{ is not empty}\}$ and $d_t = \mathbb{I}\{\text{branch node } t \text{ is an active split}\}$ be binary variables that indicate the activation status of the corresponding nodes.

Each node t is associated with a split defined by decision variables $\mathbf{a}_t \in \{0, 1\}^p$ and $b_t \in [0, 1]$, where $\mathbf{a}_t$ indicates which feature component is involved with the split, and b_t indicates the split point, assuming that all feature components have been normalized to the interval $[0, 1]$. Note that $\mathbf{a}_t$ and b_t are the decision variables.

Let $p(t)$ denote the parent node of t , and $A_L(t)$ the set of left ancestor nodes of node t , *i.e.*, the set of ancestors of t whose left branch has been followed on the path from the root to t . Similarly, we let $A_R(t)$ as the set of right ancestor nodes of t .

To ensure sufficient separation between data points in the splitting process, we define the minimum resolution for each feature dimension. Let $\varepsilon_j = \min\{x_j^{(n+1)} - x_j^{(n)} \mid x_j^{(n+1)} \neq x_j^{(n)}, n \in [N-1]\}$ is the smallest nonzero difference between observed value of feature component j , where $x_j^{(n)}$ is the n -th largest value in the j -th feature. In addition, we denote $\varepsilon_{\max} = \max_j \varepsilon_j$.

The integer programming formulation of the DDR decision tree is given as follows:

$$\begin{aligned}
\min \quad & \frac{\lambda}{N} \sum_{k \in [K]} \sum_{n \in [N]} r_{nk} \left(\frac{N}{\lambda} \|\tilde{\mathbf{z}}_n - \hat{\mathbf{z}}_k\|_2^2 + \boldsymbol{\kappa}_n^\top \mathbf{b} \right) \\
\text{s.t.} \quad & \mathbf{A}^\top \boldsymbol{\kappa}_n \geq -\mu \tilde{\mathbf{z}}_n - (1 - \mu) \hat{\mathbf{z}}_k + (1 - r_{nk})M, \quad \forall n \in [N], k \in [K] \\
& \sum_{k \in [K]} r_{nk} = 1, \quad \forall n \in [N] \\
& r_{nk} \leq g_k, \quad \forall n \in [N], l \in \mathcal{T}_K \\
& \sum_{n \in [N]} r_{nk} \geq N_{\min} g_k, \quad \forall l \in \mathcal{T}_K \\
& \mathbf{a}_m^\top \tilde{\mathbf{x}}_n \geq b_m - (1 - r_{nk}), \quad \forall l \in \mathcal{T}_K, n \in [N], m \in A_R(k) \\
& \mathbf{a}_m^\top (\tilde{\mathbf{x}}_n + \boldsymbol{\varepsilon}) \leq b_m + (1 + \varepsilon_{\max})(1 - r_{nk}), \quad \forall l \in \mathcal{T}_K, n \in [N], m \in A_L(k) \\
& \sum_{j \in [p]} a_{jt} = d_t, \quad \forall t \in \mathcal{T}_B \\
& 1 - d_t \leq b_t \leq 1, \quad \forall t \in \mathcal{T}_B \\
& d_t \leq d_{p(t)}, \quad \forall t \in \mathcal{T}_B \setminus \{1\} \\
& a_{jt}, d_t \in \{0, 1\}, \quad \forall j \in [p], t \in \mathcal{T}_B \\
& r_{nk}, g_k \in \{0, 1\}, \quad \forall n \in [N], l \in \mathcal{T}_K
\end{aligned}$$

The objective function and the first four constraints form the DDR optimization problem, while the remaining constraints are adapted from the requirements of the optimal classification tree. For more details about the tree constraints, the reader can consult ?.

The objective function $\frac{\lambda}{N} \sum_{k \in [K]} \sum_{n \in [N]} r_{nk} \left(\frac{N}{\lambda} \|\tilde{\mathbf{z}}_n - \hat{\mathbf{z}}_k\|_2^2 + \boldsymbol{\kappa}_n^\top \mathbf{b} \right)$ can be further reformulated as the following inequalities:

$$\begin{aligned}
& \frac{\lambda}{N} \sum_{k \in [K]} \sum_{n \in [N]} r_{nk} \left(\frac{N}{\lambda} \varrho_{nk} + \boldsymbol{\kappa}_n^\top \mathbf{b} \right) \\
& (\varrho_{nk} + 1)^2 \geq (\varrho_{nk} - 1)^2 + 4\varpi^2 \\
& \varpi^2 \geq \|\tilde{\mathbf{z}}_n - \hat{\mathbf{z}}_k\|_2^2
\end{aligned}$$

Moreover, the bilinear terms $r_{nk} \varrho_{nk}$ and $r_{nk} \boldsymbol{\kappa}_n$ can be linearized by introducing additional variables and applying standard linearization techniques, such as McCormick envelopes.