
Distributionally and Adversarially Robust Logistic Regression via Intersecting Wasserstein Balls

Aras Selvi¹ Eleonora Kreačić² Mohsen Ghassemi² Vamsi K. Potluru² Tucker Balch² Manuela Veloso²

¹Imperial Business School, Imperial College London

²JP Morgan AI Research

Abstract

Adversarially robust optimization (ARO) has emerged as the *de facto* standard for training models that hedge against adversarial attacks in the test stage. While these models are robust against adversarial attacks, they tend to suffer severely from overfitting. To address this issue, some successful methods replace the empirical distribution in the training stage with alternatives including (i) a worst-case distribution residing in an ambiguity set, resulting in a distributionally robust (DR) counterpart of ARO; (ii) a mixture of the empirical distribution with a distribution induced by an auxiliary (e.g., synthetic, external, out-of-domain) dataset. Inspired by the former, we study the Wasserstein DR counterpart of ARO for logistic regression and show it admits a tractable convex optimization reformulation. Adopting the latter setting, we revise the DR approach by intersecting its ambiguity set with another ambiguity set built using the auxiliary dataset, which offers a significant improvement whenever the Wasserstein distance between the data generating and auxiliary distributions can be estimated. We study the underlying optimization problem, develop efficient solution algorithms, and demonstrate that the proposed method outperforms benchmark approaches on standard datasets.

1 INTRODUCTION

Supervised learning traditionally involves access to a training dataset whose instances are assumed to be independently sampled from a true data-generating distribution (Bishop, 2006; Hastie et al., 2017). Optimizing an expected loss for the empirical distribution constructed from such a training set, also known as *empirical risk minimization* (ERM), enjoys several desirable properties in relatively generic set-

tings, including convergence to the true risk minimization problem as the number of training samples increases (Vapnik, 1999, Chapter 2). In real-world applications, however, various challenges, such as data scarcity and the existence of adversarial attacks, lead to deteriorated out-of-sample performance for models trained via ERM.

One of the key limitations of ERM, particularly as it is designed to minimize an expected loss for the empirical distribution, emerges from the finite nature of data in practice. This leads ERM to suffer from the ‘optimism bias’, also known as overfitting (Murphy, 2022), or the optimizer’s curse (DeMiguel and Nogales, 2009; Smith and Winkler, 2006), causing deteriorated out-of-sample performance. A popular approach to prevent this phenomenon, *distributionally robust optimization* (DRO; Delage and Ye 2010), optimizes the expected loss for the worst-case distribution residing within a pre-specified ambiguity set.

Another key challenge faced by ERM in practice is adversarial attacks, where an adversary perturbs the observed features during the testing or deployment phase (Szegedy et al., 2014; Goodfellow et al., 2015), also known as evasion corruption at test time (Biggio et al., 2013). For neural networks, the paradigm of *adversarial training* (AT; Madry et al. 2018) is thus designed to provide adversarial robustness by simulating such attacks in the training stage. Several successful variants of AT, specialized to different losses and attacks, have been proposed in the literature to achieve adversarial robustness without significantly reducing performance on training sets (Shafahi et al., 2019; Zhang et al., 2019; Gao et al., 2019; Pang et al., 2022). Some studies (Uesato et al., 2018; Carlini et al., 2019; Wu et al., 2020) investigate the adversarial robustness guarantees of various training algorithms, leading to a research direction focused on heuristic improvements to such models (e.g., Rade and Moosavi-Dezfooli 2022). Our work aligns with another recent line of research (Xing et al., 2022a; Bennouna et al., 2023) on *adversarially robust optimization* (ARO), which constrains ERM to guarantee an *exact*, pre-specified level of adversarial robustness while maximizing training accuracy.

Recently, it has been observed that the two aforementioned notions of robustness can be at odds, as adversarially robust (AR) models suffer from severe overfitting (*robust overfitting*; Raghunathan et al. 2019; Yu et al. 2022; Li and Spratling 2023). Indeed, it is observed that robust overfitting is even more severe than traditional overfitting (Rice et al., 2020). To this end, some works address robust overfitting by revisiting AT algorithms and adding adjustments for better generalization (Chen et al., 2020; Li and Li, 2023). In a recent work, Bennouna et al. (2023, Thm 3.2) decompose the error gap of robust overfitting into the statistical error of estimating the true data-generating distribution via the empirical distribution and an adversarial error resulting from the adversarial attacks, hence proposing the simultaneous adoption of DRO and ARO.

In this work, we study logistic regression (LR) for binary classification that is adversarially robust against ℓ_p -attacks (Crocé et al., 2020). To address robust overfitting faced by the adversarially robust LR model, we employ a DRO approach where distributional ambiguity is modeled with the type-1 Wasserstein metric. We base our work on an observation that the worst-case logistic loss under adversarial attacks can be represented as a Lipschitz continuous and convex loss function. This allows us to use existing Wasserstein DRO machinery for Lipschitz losses, and derive an *exact* reformulation of the Wasserstein DR counterpart of adversarially robust LR as a tractable convex problem.

Our main contribution lies in reducing the size of the Wasserstein ambiguity set in the DRO problem mentioned above, in order to create a less conservative problem while preserving the same distributional robustness guarantees. To accomplish this, we draw inspiration from recent work on ARO that leverages auxiliary datasets (*e.g.*, [Gowal et al. 2021](#); [Xing et al. 2022a](#)) and revise our DRO problem by intersecting its ambiguity set with another ambiguity set constructed using an auxiliary dataset. Examples of auxiliary data include synthetic data generated from a generative model (*e.g.*, privacy-preserving data release), data in the presence of distributional shifts (*e.g.*, different time periods/regions), noisy data (*e.g.*, measurement errors), or out-of-domain data (*e.g.*, different source); any auxiliary dataset is viable as long as its instances are sampled independently from an underlying data-generating distribution whose Wasserstein distance to the true data-generating distribution is known or can be estimated. [Figure 1](#) illustrates our framework.

The paper unfolds as follows. In Section 2, we review related literature on DRO and ARO, with a focus on their interactions. We examine the use of auxiliary data in ARO and the intersection of Wasserstein balls in DRO. We discuss open questions for LR to motivate our loss function choice in this work. Section 3 gives preliminaries on ERM, ARO, and type-1 Wasserstein DRO. In Section 4, we discuss that the adversarial logistic loss can be reformulated as a Lipschitz convex function, enabling the use of Wasserstein

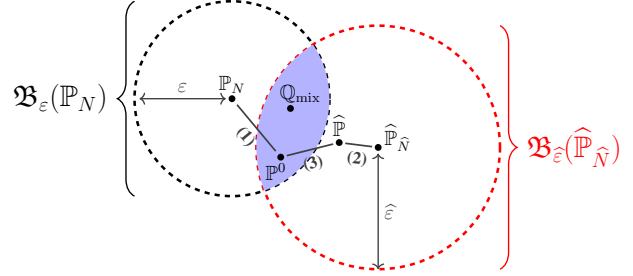


Figure 1: Traditional ARO optimizes the expected adversarial loss over the empirical distribution $\mathbb{P}_N$ constructed from N i.i.d. samples of the (unknown) true data-generating distribution $\mathbb{P}^0$. Replacing $\mathbb{P}_N$ with a worst-case distribution in the ball $\mathfrak{B}_\varepsilon(\mathbb{P}_N)$ gives us its DR counterpart. To reduce the size of this ball, we intersect it with another ball $\mathfrak{B}_\varepsilon(\hat{\mathbb{P}}_{\hat{N}})$ while ensuring $\mathbb{P}^0$ is still included with high confidence. The latter ball is centered at an empirical distribution $\hat{\mathbb{P}}_{\hat{N}}$ constructed from $\hat{N}$ i.i.d. samples of some auxiliary distribution $\hat{\mathbb{P}}$. Recent works using auxiliary data in ARO propose optimizing the expected adversarial loss over a mixture $\mathbb{Q}_{\text{mix}}$ of $\mathbb{P}_N$ and $\hat{\mathbb{P}}_{\hat{N}}$; we show that this distribution resides in $\mathfrak{B}_\varepsilon(\mathbb{P}_N) \cap \mathfrak{B}_\varepsilon(\hat{\mathbb{P}}_{\hat{N}})$ under some conditions.

DRO machinery for Lipschitz losses. Our main contribution (cf. Figure 1) is in Section 5, where we provide an explicit reformulation of the distributionally and adversarially robust LR problem over the intersection of two Wasserstein balls, prove the NP-hardness of this problem, and derive a convex relaxation of it. Our work is mainly on *optimization* where we focus on how to solve the underlying problems upon cross-validating Wasserstein ball radii, however, in Section 6 we discuss some preliminary statistical approaches to set such radii. We close the paper with numerical experiments on standard benchmark datasets in Section 7. We borrow the standard notation in DR machine learning, which is elaborated on in our Appendices.

2 RELATED WORK

Auxiliary data in ARO. The use of auxiliary data appears in the ARO literature. In particular, it is shown that additional unlabeled data sampled from the same (Carmon et al., 2019; Xing et al., 2022b) or different (Deng et al., 2021) data-generating distributions could provide adversarial robustness. Sehwal et al. (2022) show that adversarial robustness can be certified even when it is provided for a synthetic dataset as long as the distance between its generator and the true data-generating distribution can be quantified. Goyal et al. (2021) and Xing et al. (2022a) propose optimizing a weighted combination of ARO over empirical and synthetic datasets. We show that the latter approach can be recovered by our model.

DRO-ARO interactions. In this work, we optimize ARO

against worst-case data-generating distributions residing in an ambiguity set, where the type-1 Wasserstein metric is used for distances since it is arguably the most common choice in machine learning (ML) with Lipschitz losses (Shafieezadeh-Abadeh et al., 2019; Gao, 2023). In the literature, it is shown that standard ARO is equivalent to the DRO of the original loss function with a type- ∞ Wasserstein metric (Staib and Jegelka, 2017; Khim and Loh, 2018; Pydi and Jog, 2021; Regniet et al., 2022; Frank and Niles-Weed, 2024). In other words, in the absence of adversarial attacks, training models adversarially with artificial attacks provide some distributional robustness. Hence, our DR ARO approach can be interpreted as optimizing the logistic loss over the worst-case distribution whose 1-Wasserstein distance is bounded by a pre-specified radius from at least one distribution residing in an ∞ -Wasserstein ball around the empirical distribution. Conversely, Sinha et al. (2018) discuss that while DRO over Wasserstein balls is intractable for generic losses (e.g., neural networks), its Lagrange relaxation resembles ARO and thus ARO yields a certain degree of (relaxed) distributional robustness (Wu et al., 2020; Bui et al., 2022; Phan et al., 2023). Such literature suggests that, when there is no concern about the statistical errors caused by using empirical distributions (e.g., in very high-data regimes), one can train DR models to obtain adversarial robustness guarantees. However, as discussed by Bennouna and Van Parys (2022), when statistical errors exist, then we need to be simultaneously robust against adversarial attacks and statistical errors. To the best of our knowledge, there have not been works optimizing a pre-specified level of type-1 Wasserstein distributional robustness (that hedges against overfitting, Kuhn et al. 2019) and adversarial robustness (that hedges against adversarial attacks, Goodfellow et al. 2015) *simultaneously*. To our knowledge, the only approach that considers the exact DR counterpart of ARO is proposed by Bennouna et al. (2023), who model distributional ambiguity with φ -divergences for neural networks.

Intersecting ambiguity sets in DRO. Recent work started to explore the intersection of ambiguity sets for different contexts (Awasthi et al., 2022; Wang et al., 2024) or different metrics (Tanoumand et al., 2023). Our idea of intersecting Wasserstein balls is originated from the “Surround, then Intersect” strategy (Taskesen et al., 2021, §5.2) to train linear regression under sequential domain adaptation in a non-adversarial setting (see the work of Shafahi et al. 2020 and Song et al. 2019 for robustness in domain adaptation/transfer learning). The aforementioned work focuses on the squared loss function with an ambiguity set using the Wasserstein metric developed for the first and second distributional moments. In a recent study, Rychener et al. (2024) generalize most of the previous results and prove that DRO problems over the intersection of two Wasserstein balls admit tractable convex reformulations whenever the loss function is the maximum of concave functions. They also discuss why distributions lying in the intersection of

two Wasserstein balls are more natural candidates for the unknown true distribution than those that are Wasserstein barycenters or mixture distributions of the empirical and auxiliary distributions (referred to as heterogeneous data sources; see Example 1, Proposition 2, and Corollary 1).

Logistic loss in DRO and ARO. Our choice of LR aligns with the current directions and open questions in the related literature. In the DRO literature, even in the absence of adversarial attacks, the aforementioned work of Taskesen et al. (2021) on the intersection of Wasserstein ambiguity sets is restricted to linear regression. The authors show that this problem admits a tractable convex optimization reformulation, and their proof relies on the properties of the squared loss. Similarly, Rychener et al. (2024) discuss that the logistic loss fails to satisfy the piece-wise concavity assumption and is inherently difficult to optimize over the intersection of Wasserstein balls. We contribute to the DRO literature for adversarial and non-adversarial settings because we show that such a problem would be NP-hard for the logistic loss even without adversarial attacks, and develop specialized approximation techniques. Our problem recovers DR LR (Shafieezadeh-Abadeh et al., 2015; Selvi et al., 2022) as a special case in the absence of adversarial attacks and auxiliary data. Answering theoretical challenges posed by logistic regression has been useful in answering more general questions in the DRO literature, such as DR LR (Shafieezadeh-Abadeh et al., 2015) leading to DR ML (Shafieezadeh-Abadeh et al., 2019) and mixed-feature DR LR (Selvi et al., 2022) leading to mixed-feature DR Lipschitz ML (Belbasi et al., 2023). Finally, in the (non-DR) ARO literature, there are recent theory developments on understanding the effect of auxiliary data (e.g., Xing et al. 2022a) specifically for squared and logistic loss functions.

Single-step adversarial training and single-shot ARO.

Our work proposes a single-shot convex optimization procedure to train logistic models that are both adversarially and distributionally robust. Although the terminology may resemble the recent work on *single-step adversarial training* for neural networks (Wong et al., 2020; Lin et al., 2023, 2024), the two approaches operate differently. Single-step adversarial training generates adversarial perturbations using a single gradient computation at each iteration of iterative model training and improves robustness through updates, with performance typically assessed at intermediate checkpoints. In contrast, our method solves a convex optimization problem once to obtain a model that satisfies both forms of robustness by design. To enable tractability, this approach leverages the convexity and Lipschitz continuity of the loss function under adversarial attacks, which hold for the logistic loss function. While single-step adversarial training applies broadly to general classes of models such as neural networks, our framework offers a complementary, optimization-based perspective in the logistic regression model, where structural properties can be fully exploited.

3 PRELIMINARIES

We consider a binary classification problem where an instance is modeled as $(x, y) \in \Xi := \mathbb{R}^n \times \{-1, +1\}$ and the labels depend on the features via $\text{Prob}[y \mid x] = [1 + \exp(-y \cdot \beta^\top x)]^{-1}$ for some $\beta \in \mathbb{R}^n$; its associated loss is the *logloss* $\ell_\beta(x, y) := \log(1 + \exp(-y \cdot \beta^\top x))$.

Empirical risk minimization. Let $\mathcal{P}(\Xi)$ denote the set of distributions supported on Ξ and $\mathbb{P}^0 \in \mathcal{P}(\Xi)$ denote the true data-generating distribution. One wants to minimize the expected logloss over $\mathbb{P}^0$, that is

$$\inf_{\beta \in \mathbb{R}^n} \mathbb{E}_{\mathbb{P}^0}[\ell_\beta(x, y)]. \quad (\text{RM})$$

In practice, $\mathbb{P}^0$ is hardly ever known, and one resorts to the empirical distribution $\mathbb{P}_N = \frac{1}{N} \sum_{i \in [N]} \delta_{\xi^i}$ where $\xi^i = (x^i, y^i)$, $i \in [N]$, are i.i.d. samples from $\mathbb{P}^0$ and δ_ξ denotes the Dirac distribution supported on ξ . The empirical risk minimization (ERM) problem is thus

$$\inf_{\beta \in \mathbb{R}^n} \mathbb{E}_{\mathbb{P}_N}[\ell_\beta(x, y)]. \quad (\text{ERM})$$

Distributionally robust optimization. To be able to define a distance between distributions, we first define the following feature-label metric on Ξ .

Definition 1. The distance between instances $\xi = (x, y) \in \Xi$ and $\xi' = (x', y') \in \Xi$ for $\kappa \geq 0$ and $q \geq 1$ is

$$d(\xi, \xi') = \|x - x'\|_q + \kappa \cdot \mathbb{1}[y \neq y'].$$

Using this metric, we define the Wasserstein distance.

Definition 2. The type-1 Wasserstein distance between distributions $\mathbb{Q}, \mathbb{Q}' \in \mathcal{P}(\Xi)$ is defined as

$$W(\mathbb{Q}, \mathbb{Q}') = \inf_{\Pi \in \mathcal{C}(\mathbb{Q}, \mathbb{Q}')} \left\{ \int_{\Xi \times \Xi} d(\xi, \xi') \Pi(d\xi, d\xi') \right\},$$

where $\mathcal{C}(\mathbb{Q}, \mathbb{Q}')$ is the set of couplings of $\mathbb{Q}$ and $\mathbb{Q}'$.

In finite-data settings, the distance between the true data-generating distribution and the empirical distribution is upper-bounded by some $\epsilon > 0$. The Wasserstein DRO problem is thus defined as

$$\inf_{\beta \in \mathbb{R}^n} \sup_{\mathbb{Q} \in \mathfrak{B}_\epsilon(\mathbb{P}_N)} \mathbb{E}_{\mathbb{Q}}[\ell_\beta(x, y)], \quad (\text{DRO})$$

where $\mathfrak{B}_\epsilon(\mathbb{P}) := \{\mathbb{Q} \in \mathcal{P}(\Xi) : W(\mathbb{Q}, \mathbb{P}) \leq \epsilon\}$ denotes the Wasserstein ball centered at $\mathbb{P} \in \mathcal{P}(\Xi)$ with radius ϵ . We refer to Mohajerin Esfahani and Kuhn (2018) and Kuhn et al. (2019) for the properties of DRO and estimating ϵ .

Adversarially robust optimization. The goal of adversarial robustness is to provide robustness against adversarial attacks (Goodfellow et al., 2015). An adversarial attack, in the widely studied ℓ_p -noise setting (Croce et al., 2020), perturbs the features of the test instances (x, y) by adding additive noise z to x . The adversary chooses the noise vector z , subject to $\|z\|_p \leq \alpha$, so as to maximize the loss $\ell_\beta(x + z, y)$ associated with this perturbed test instance. Therefore, ARO solves the following optimization problem in the training stage to hedge against adversarial perturbations at the test stage:

$$\inf_{\beta \in \mathbb{R}^n} \mathbb{E}_{\mathbb{P}_N} \left[\sup_{z: \|z\|_p \leq \alpha} \{\ell_\beta(x + z, y)\} \right]. \quad (\text{ARO})$$

ARO reduces to ERM when $\alpha = 0$. Note that ARO is identical to feature robust training (Bertsimas et al., 2019) which is not motivated by adversarial attacks, but by the presence of noisy observations in the training set (Ben-Tal et al., 2009; Gorissen et al., 2015).

DRO-ARO connection. A connection between ARO and DRO is noted in the literature (Staub and Jegelka 2017, Proposition 3.1, Khim and Loh 2018, Lemma 22, Pydi and Jog 2021, Lemma 5.1, Regniet et al. 2022, Proposition 2.1, Frank and Niles-Weed 2024, Lemma 3, and Bennouna et al. 2023, §3). Namely, problem ARO is equivalent to a DRO problem

$$\inf_{\beta \in \mathbb{R}^n} \sup_{\mathbb{Q} \in \mathfrak{B}_\alpha^\infty(\mathbb{P}_N)} \mathbb{E}_{\mathbb{Q}}[\ell_\beta(x, y)], \quad (1)$$

where the ambiguity set $\mathfrak{B}_\alpha^\infty(\mathbb{P}_N)$ is a type- ∞ Wasserstein ball (Givens and Shortt, 1984) with radius α . Hence, in non-adversarial settings, ARO provides robustness with respect to the type- ∞ Wasserstein distance. In the case of adversarial attacks, it suffers from robust overfitting as discussed earlier. To address this issue, one straightforward approach is to revisit (1) and replace α with some $\alpha' > \alpha$. This approach, however, does not provide improvements for the out-of-sample performance since (i) the type- ∞ Wasserstein distance employed in problem (1) uses a metric on the feature space, ignoring labels; (ii) type- ∞ Wasserstein distances do not provide strong out-of-sample performances in ML (unlike, e.g., the type-1 Wasserstein distance) since the required radii to provide meaningful robustness guarantees are typically too large (Bennouna and Van Parys, 2022, §1.2.2, and references therein). We thus study the type-1 Wasserstein counterpart of ARO, which we initiate in the next section.

4 DISTRIBUTIONALLY AND ADVERSARIALLY ROBUST LR

Here we derive the Wasserstein DR counterpart of ARO that will set the ground for our main result in the next section. We impose the following assumption.

Assumption 1. We are given a finite $\varepsilon > 0$ value satisfying $W(\mathbb{P}^0, \mathbb{P}_N) \leq \varepsilon$.

The assumption implies that we know an $\varepsilon > 0$ value satisfying $\mathbb{P}^0 \in \mathfrak{B}_\varepsilon(\mathbb{P}_N)$. Typically, however, ε is either estimated through cross-validation or finite sample statistics, with the assumption then regarded as holding with high confidence (see §6 for a review of related results we can borrow). The distributionally and adversarially robust LR problem is thus:

$$\inf_{\beta \in \mathbb{R}^n} \sup_{Q \in \mathfrak{B}_\varepsilon(\mathbb{P}_N)} \mathbb{E}_Q[\sup_{z: \|z\|_p \leq \alpha} \{\ell_\beta(x + z, y)\}]. \quad (\text{DR-ARO})$$

By employing a simple duality trick for the inner sup-problem, as commonly applied in robust optimization (Ben-Tal et al., 2009; Bertsimas and Den Hertog, 2022), we can represent **DR-ARO** as a standard non-adversarial DRO problem with an updated loss function, which we name the *adversarial loss*.

Observation 1. Problem **DR-ARO** is equivalent to

$$\inf_{\beta \in \mathbb{R}^n} \sup_{Q \in \mathfrak{B}_\varepsilon(\mathbb{P}_N)} \mathbb{E}_Q[\ell_\beta^\alpha(x, y)],$$

where the adversarial loss ℓ_β^α is defined as

$$\ell_\beta^\alpha(x, y) := \log(1 + \exp(-y \cdot \beta^\top x + \alpha \cdot \|\beta\|_{p^*})),$$

for p^* satisfying $1/p + 1/p^* = 1$. The univariate representation $L^\alpha(z) := \log(1 + \exp(-z + \alpha \cdot \|\beta\|_{p^*}))$ of ℓ_β^α is convex and has a Lipschitz modulus of 1.

As a corollary of Observation 1, we can directly employ the techniques proposed by Shafieezadeh-Abadeh et al. (2019) to dualize the inner sup-problem of **DR-ARO** and obtain a tractable reformulation.

Corollary 1. Problem **DR-ARO** admits the following tractable convex optimization reformulation:

$$\begin{aligned} \inf_{\beta, \lambda, \mathbf{s}} \quad & \varepsilon \lambda + \frac{1}{N} \sum_{i=1}^N s_i \\ \text{s.t.} \quad & \ell_\beta^\alpha(x^i, y^i) \leq s_i \quad \forall i \in [N] \\ & \ell_\beta^\alpha(x^i, -y^i) - \lambda \kappa \leq s_i \quad \forall i \in [N] \\ & \|\beta\|_{q^*} \leq \lambda \\ & \beta \in \mathbb{R}^n, \lambda \geq 0, \mathbf{s} \in \mathbb{R}_+^N, \end{aligned}$$

for q^* satisfying $1/q + 1/q^* = 1$.

The constraints of this problem are exponential cone representable (derivation is in the appendices) and for $q \in \{1, 2, \infty\}$, the yielding problem can be solved with the exponential cone solver of MOSEK (MOSEK ApS, 2023) in polynomial time (Nesterov, 2018).

5 MAIN RESULT

In §4 we discussed the traditional DRO setting where we have access to an empirical distribution $\mathbb{P}_N$ constructed from N i.i.d. samples of the true data-generating distribution $\mathbb{P}^0$, and we are given (or we estimate) some ε so that $\mathbb{P}^0 \in \mathfrak{B}_\varepsilon(\mathbb{P}_N)$. Recently in DRO literature, it became a key focus to study the case where we have access to an additional auxiliary empirical distribution $\hat{\mathbb{P}}_{\hat{N}}$ constructed from $\hat{N}$ i.i.d. samples $\hat{\xi}^j = (\hat{x}^j, \hat{y}^j)$, $j \in [\hat{N}]$, of some other distribution $\hat{\mathbb{P}}$; given the increasing availability of useful auxiliary data in the ARO domain, we explore this direction here. We start with the following assumption.

Assumption 2. We are given finite $\varepsilon, \hat{\varepsilon} > 0$ values satisfying $W(\mathbb{P}^0, \mathbb{P}_N) \leq \varepsilon$ and $W(\mathbb{P}^0, \hat{\mathbb{P}}_{\hat{N}}) \leq \hat{\varepsilon}$.

The assumption implies that we know $\varepsilon, \hat{\varepsilon} > 0$ values satisfying $\mathbb{P}^0 \in \mathfrak{B}_\varepsilon(\mathbb{P}_N) \cap \mathfrak{B}_{\hat{\varepsilon}}(\hat{\mathbb{P}}_{\hat{N}})$. In practice, this assumption is ensured to hold with high confidence by estimating the ε and $\hat{\varepsilon}$ values; methods across various domains which we can adopt are reviewed in §6. We want to optimize the adversarial loss over the intersection $\mathfrak{B}_\varepsilon(\mathbb{P}_N) \cap \mathfrak{B}_{\hat{\varepsilon}}(\hat{\mathbb{P}}_{\hat{N}})$:

$$\inf_{\beta \in \mathbb{R}^n} \sup_{Q \in \mathfrak{B}_\varepsilon(\mathbb{P}_N) \cap \mathfrak{B}_{\hat{\varepsilon}}(\hat{\mathbb{P}}_{\hat{N}})} \mathbb{E}_Q[\ell_\beta^\alpha(x, y)]. \quad (\text{Inter-ARO})$$

Note that Assumption 2 implies that ε and $\hat{\varepsilon}$ guarantee that $\mathfrak{B}_\varepsilon(\mathbb{P}_N) \cap \mathfrak{B}_{\hat{\varepsilon}}(\hat{\mathbb{P}}_{\hat{N}})$ is nonempty, and problem **Inter-ARO** is thus feasible. This problem is expected to outperform **DR-ARO** as the ambiguity set is smaller while still including $\mathbb{P}^0$. However, problem **Inter-ARO** is challenging to solve even in the absence of adversarial attacks ($\alpha = 0$) as we reviewed in §2. To address this challenge, we first reformulate **Inter-ARO** as a semi-infinite optimization problem with finitely many variables.

Proposition 1. **Inter-ARO** is equivalent to:

$$\begin{aligned} \inf_{\beta, \lambda, \hat{\lambda}, \mathbf{s}, \hat{\mathbf{s}}} \quad & \varepsilon \lambda + \hat{\varepsilon} \hat{\lambda} + \frac{1}{N} \sum_{i=1}^N s_i + \frac{1}{\hat{N}} \sum_{j=1}^{\hat{N}} \hat{s}_j \\ \text{s.t.} \quad & \sup_{x \in \mathbb{R}^n} \{\ell_\beta^\alpha(x, l) - \lambda \|x^i - x\|_q - \hat{\lambda} \|\hat{x}^j - x\|_q\} \\ & \leq s_i + \frac{\kappa(1 - ly^i)}{2} \lambda + \hat{s}_j + \frac{\kappa(1 - l\hat{y}^j)}{2} \hat{\lambda} \\ & \quad \forall (i, j, l) \in [N] \times [\hat{N}] \times \{-1, 1\} \\ & \beta \in \mathbb{R}^n, \lambda \geq 0, \hat{\lambda} \geq 0, \mathbf{s} \in \mathbb{R}_+^N, \hat{\mathbf{s}} \in \mathbb{R}_+^{\hat{N}}. \end{aligned}$$

Even though this problem recovers the tractable problem **DR-ARO** as $\hat{\varepsilon} \rightarrow \infty$, it is NP-hard in the finite radius settings. We reformulate **Inter-ARO** as an adjustable robust optimization problem (Ben-Tal et al., 2004; Yanikoğlu et al., 2019), and borrow tools from this literature to obtain the following result.

Proposition 2. **Inter-ARO** is equivalent to an adjustable RO problem with $\mathcal{O}(N \cdot \hat{N})$ two-stage robust constraints, which is NP-hard even when $N = \hat{N} = 1$.

The adjustable RO literature has developed a rich arsenal of relaxations that can be leveraged for **Inter-ARO**. We adopt the ‘static relaxation technique’ (Bertsimas et al., 2015) to restrict the feasible region of **Inter-ARO** and obtain a tractable approximation.

Theorem 1 (main). *The following convex optimization problem is a feasible relaxation of **Inter-ARO**:*

$$\begin{aligned}
& \inf_{\substack{\beta, \lambda, \hat{\lambda}, \mathbf{s}, \hat{\mathbf{s}} \\ \mathbf{z}_{ij}^+, \mathbf{z}_{ij}^- \\ \text{s.t.}}} \quad \varepsilon \lambda + \hat{\varepsilon} \hat{\lambda} + \frac{1}{N} \sum_{i=1}^N s_i + \frac{1}{\hat{N}} \sum_{j=1}^{\hat{N}} \hat{s}_j \\
& \quad L^\alpha(l \cdot \beta^\top \mathbf{x}^i + \mathbf{z}_{ij}^{l\top} (\hat{\mathbf{x}}^j - \mathbf{x}^i)) \\
& \quad \leq s_i + \frac{\kappa(1 - ly^i)}{2} \lambda + \hat{s}_j + \frac{\kappa(1 - l\hat{y}^j)}{2} \hat{\lambda} \\
& \quad \quad \forall (i, j, l) \in [N] \times [\hat{N}] \times \{-1, 1\} \\
& \quad \|\mathbf{l}\beta - \mathbf{z}_{ij}^l\|_{q^*} \leq \lambda \\
& \quad \quad \forall (i, j, l) \in [N] \times [\hat{N}] \times \{-1, 1\} \\
& \quad \|\mathbf{z}_{ij}^l\|_{q^*} \leq \hat{\lambda} \\
& \quad \quad \forall (i, j, l) \in [N] \times [\hat{N}] \times \{-1, 1\} \\
& \quad \beta \in \mathbb{R}^n, \lambda \geq 0, \hat{\lambda} \geq 0, \mathbf{s} \in \mathbb{R}_+^N, \hat{\mathbf{s}} \in \mathbb{R}_+^{\hat{N}} \\
& \quad \mathbf{z}_{ij}^l \in \mathbb{R}^n, (i, j, l) \in [N] \times [\hat{N}] \times \{-1, 1\}. \\
& \quad \quad \quad \text{(Inter-ARO*)}
\end{aligned}$$

Similarly to **DR-ARO**, the constraints of **Inter-ARO*** are exponential cone representable (cf. appendices).

Recall that for $\hat{\varepsilon}$ large enough, **Inter-ARO** reduces to **DR-ARO**. The following corollary shows that, despite **Inter-ARO*** being a relaxation of **Inter-ARO**, a similar property holds. That is, “not learning anything from auxiliary data” remains feasible: the static relaxation does not force learning from $\hat{\mathbb{P}}_{\hat{N}}$, and it learns from auxiliary data only if the objective improves.

Corollary 2. *Feasibility of ignoring auxiliary data: Any feasible solution $(\beta, \lambda, \mathbf{s})$ of **DR-ARO** can be used to recover a feasible solution $(\beta, \lambda, \hat{\lambda}, \mathbf{s}, \hat{\mathbf{s}}, \mathbf{z}_{ij}^+, \mathbf{z}_{ij}^-)$ for **Inter-ARO*** with $\hat{\lambda} = 0$, $\hat{\mathbf{s}} = \mathbf{0}$, and $\mathbf{z}_{ij}^+ = \mathbf{z}_{ij}^- = \mathbf{0}$. Convergence to **Inter-ARO**: The optimal value of **Inter-ARO*** converges to the optimal value of **Inter-ARO**, with the same set of β solutions, as $\hat{\varepsilon} \rightarrow \infty$.*

In light of Corollary 2, Appendix D.1 discusses that some simulations in our numerical experiments chose not to incorporate the auxiliary data by setting a sufficiently large $\hat{\varepsilon}$. We close the section by discussing how **Inter-ARO** can recover some problems in the DRO and ARO literature. Firstly, recall that **Inter-ARO** can ignore the auxiliary data once $\hat{\varepsilon}$ is set large enough, reducing this problem to **DR-ARO**. Moreover, notice that $\alpha = 0$ reduces ℓ_β^α to ℓ_β , hence for $\alpha = 0$ and $\hat{\varepsilon} = \infty$ **Inter-ARO** recovers the Wasserstein LR model of Shafieezadeh-Abadeh et al. (2015). We next

relate **Inter-ARO** to the problems in the ARO literature that use auxiliary data. The works in this literature (Gowal et al., 2021; Xing et al., 2022a) solve the following

$$\inf_{\beta \in \mathbb{R}^n} \frac{1}{N + w\hat{N}} \left[\sum_{i \in [N]} \sup_{\mathbf{z}^i \in \mathcal{B}_p(\alpha)} \{\ell_\beta(\mathbf{x}^i + \mathbf{z}^i, y^i)\} + w \sum_{j \in [\hat{N}]} \sup_{\mathbf{z}^j \in \mathcal{B}_p(\alpha)} \{\ell_\beta(\hat{\mathbf{x}}^j + \mathbf{z}^j, \hat{y}^j)\} \right], \quad (2)$$

for some $w > 0$, where $\mathcal{B}_p(\alpha) := \{\mathbf{z} \in \mathbb{R}^n : \|\mathbf{z}\|_p \leq \alpha\}$. We observe that (2) resembles a variant of **ARO** that replaces the empirical distribution $\mathbb{P}_N$ with its mixture with $\hat{\mathbb{P}}_{\hat{N}}$:

Observation 2. *Problem (2) is equivalent to*

$$\inf_{\beta \in \mathbb{R}^n} \mathbb{E}_{\mathbb{Q}_{\text{mix}}} [\ell_\beta(\mathbf{x}, y)] \quad (3)$$

where $\mathbb{Q}_{\text{mix}} := \lambda \cdot \mathbb{P}_N + (1 - \lambda) \cdot \hat{\mathbb{P}}_{\hat{N}}$ for $\lambda = \frac{N}{N + w\hat{N}}$.

We give a condition on ε and $\hat{\varepsilon}$ to guarantee that the mixture distribution introduced in Proposition 2 lives in $\mathfrak{B}_\varepsilon(\mathbb{P}_N) \cap \mathfrak{B}_\varepsilon(\hat{\mathbb{P}}_{\hat{N}})$, that is, the distribution $\mathbb{Q}_{\text{mix}}$ will be feasible in the sup problem of **Inter-ARO**.

Proposition 3. *For any $\lambda \in (0, 1)$ and $\mathbb{Q}_{\text{mix}} = \lambda \cdot \mathbb{P}_N + (1 - \lambda) \cdot \hat{\mathbb{P}}_{\hat{N}}$, we have $\mathbb{Q}_{\text{mix}} \in \mathfrak{B}_\varepsilon(\mathbb{P}_N) \cap \mathfrak{B}_\varepsilon(\hat{\mathbb{P}}_{\hat{N}})$ whenever $\varepsilon + \hat{\varepsilon} \geq W(\mathbb{P}_N, \hat{\mathbb{P}}_{\hat{N}})$ and $\frac{\hat{\varepsilon}}{\varepsilon} = \frac{\lambda}{1 - \lambda}$.*

For $\lambda = \frac{N}{N + \hat{N}}$, if the intersection $\mathfrak{B}_\varepsilon(\mathbb{P}_N) \cap \mathfrak{B}_{\hat{\varepsilon}}(\hat{\mathbb{P}}_{\hat{N}})$ is nonempty, Proposition 3 implies that a sufficient condition for this intersection to include $\mathbb{Q}_{\text{mix}}$ is $\hat{\varepsilon}/\varepsilon = N/\hat{N}$, which is intuitive since the radii of Wasserstein ambiguity sets are chosen inversely proportional to the number of samples (Kuhn et al., 2019, Theorem 18).

6 SETTING WASSERSTEIN RADII

Thus far, we have assumed knowledge of DRO ball radii ε and $\hat{\varepsilon}$ that satisfy Assumptions 1 and 2. In this section, we employ Wasserstein finite-sample statistics techniques to estimate these values.

Setting ε for **DR-ARO.** In the following theorem, we present tight characterizations for ε so that the ball $\mathfrak{B}_\varepsilon(\mathbb{P}_N)$ includes the true distribution $\mathbb{P}^0$ with arbitrarily high confidence. We show that for an ε chosen in such a manner, **DR-ARO** is well-defined. The full description of this result is available in our appendices.

Theorem 2 (abridged collection of results from Fournier and Guillin 2015; Kuhn et al. 2019; Yue et al. 2022). *For light-tailed distribution $\mathbb{P}^0$ and $\varepsilon \geq \mathcal{O}(\frac{\log(\eta^{-1})}{N})^{1/n}$ for $\eta \in (0, 1)$, we have: (i) $\mathbb{P}^0 \in \mathfrak{B}_\varepsilon(\mathbb{P}_N)$ with $1 - \eta$ confidence; (ii) **DR-ARO** overestimates the expected loss for $\mathbb{P}^0$ with $1 - \eta$ confidence; (iii) **DR-ARO** is asymptotically consistent $\mathbb{P}^0$ -a.s.; (iv) worst-case distributions for optimal solutions of **DR-ARO** are supported on at most $N + 1$ outcomes.*

We next derive an analogous result for **Inter-ARO**.

Choosing ϵ and $\hat{\epsilon}$ in **Inter-ARO.** Recall that **Inter-ARO** revises **DR-ARO** by intersecting $\mathcal{B}_\epsilon(\mathbb{P}_N)$ with $\mathcal{B}_{\hat{\epsilon}}(\hat{\mathbb{P}}_{\hat{N}})$. We need a nonempty intersection for **Inter-ARO** to be well-defined. A necessary and sufficient condition follows from the triangle inequality $\epsilon + \hat{\epsilon} \geq W(\mathbb{P}_N, \hat{\mathbb{P}}_{\hat{N}})$, where $W(\mathbb{P}_N, \hat{\mathbb{P}}_{\hat{N}})$ can be computed with linear optimization as both distributions are discrete. We also want this intersection to include $\mathbb{P}^0$ with high confidence, in order to satisfy Assumption 2. We next provide a tight characterization for such $\epsilon, \hat{\epsilon}$. The full description of this result is available in our appendices.

Theorem 3 (abridged). *For light-tailed $\mathbb{P}^0$ and $\hat{\mathbb{P}}$, if $\epsilon \geq \mathcal{O}(\frac{\log(\eta_1^{-1})}{N})^{1/n}$ and $\hat{\epsilon} \geq W(\mathbb{P}^0, \hat{\mathbb{P}}) + \mathcal{O}(\frac{\log(\eta_2^{-1})}{\hat{N}})^{1/n}$ for $\eta_1, \eta_2 \in (0, 1)$ with $\eta := \eta_1 + \eta_2 < 1$, we have: (i) $\mathbb{P}^0 \in \mathcal{B}_\epsilon(\mathbb{P}_N) \cap \mathcal{B}_{\hat{\epsilon}}(\hat{\mathbb{P}}_{\hat{N}})$ with $1 - \eta$ confidence; (ii) **Inter-ARO** overestimates true loss with $1 - \eta$ confidence.*

Remark 1. **Inter-ARO** is not asymptotically consistent, given that $\hat{N} \rightarrow \infty$ will let $\hat{\epsilon} \rightarrow W(\mathbb{P}^0, \hat{\mathbb{P}})$ due to the non-zero constant distance between the true distribution $\mathbb{P}^0$ and the auxiliary distribution $\hat{\mathbb{P}}$. **Inter-ARO** is thus not useful in asymptotic data regimes.

Remark 2. The assumption that true data-generating distributions are light-tailed is satisfied when Ξ is compact, and it is a common assumption for even simple sample average approximation techniques (Mohajerin Esfahani and Kuhn, 2018, Assumption 3.3).

Knowledge of $W(\mathbb{P}^0, \hat{\mathbb{P}})$. In Theorem 3, we use $W(\mathbb{P}^0, \hat{\mathbb{P}})$ explicitly. This distance, however, is typically unknown, and a common approach is to cross-validate it¹. This would be applicable in our setting thanks to Corollary 2, because the relaxation **Inter-ARO*** does not force learning from the auxiliary data unless it is useful, that is, one can seek evidence for the usefulness of the auxiliary data via cross-validation. Moreover, there are several domains where $W(\mathbb{P}^0, \hat{\mathbb{P}})$ is known exactly. For some special cases, we can use direct domain knowledge (e.g., the ‘‘Uber vs Lyft’’ example of Taske-sen et al. 2021). A recent example comes from learning from multi-source data, where $\mathbb{P}^0$ is named the target distribution and $\hat{\mathbb{P}}$ is the source distribution (Rychener et al., 2024, §1). Another domain is private data release, where a data holder shares a subset of opt-in data to form $\mathbb{P}_N$, and generates a privacy-preserving synthetic dataset from the rest. The (privately generated) synthetic distribution has a known nonzero Wasserstein distance from the true data-generating distribution (Dwork and Roth, 2014; Ullman and Vadhan, 2020). See (Rychener et al., 2024, §5) for more scenarios that enable quantifying $W(\mathbb{P}^0, \hat{\mathbb{P}})$. Alternatively,

one can directly rely on $W(\mathbb{P}_N, \hat{\mathbb{P}})$ if it is known, especially when synthetic data generators are trained on the empirical dataset. By employing Wasserstein GANs, which minimize the Wasserstein-1 distance, the distance between the generated distribution and the training distribution is minimized. This ensures that the synthetic distribution remains within a radius of the training distribution (Arjovsky et al., 2017).

7 EXPERIMENTS

We conduct a series of experiments, each having a different source of auxiliary data, to test the proposed DR ARO models. We use the following abbreviations, where ‘solution’ refers to the optimal β to make decisions:

- ERM: Solution of problem **ERM** (i.e., naïve LR);
- ARO: Solution of problem **ARO** (i.e., adversarially robust LR);
- ARO+Aux: Solution of problem (2) (i.e., replacing the empirical distribution of **ARO** with its mixture with auxiliary data);
- DRO+ARO: Solution of **DR-ARO** (i.e., the Wasserstein DR counterpart of **ARO**);
- DRO+ARO+Aux: Solution of **Inter-ARO*** (i.e., relaxation of **Inter-ARO** that intersects the ambiguity set of **DR-ARO** with an auxiliary Wasserstein ball);

All parameters are 5-fold cross-validated from various grids. The Wasserstein radii use the grid $\{10^{-6}, 10^{-5}, 10^{-4}, 10^{-3}, 10^{-2}, 10^{-1}, 0, 1, 2, 5, 10\}$, which is sufficient to ensure that the rule-of-thumb $\epsilon = \mathcal{O}(1/\sqrt{N})$ is included around the center of this grid for all experiments conducted. To ensure that the intersections of Wasserstein balls are nonempty, we compute $W(\mathbb{P}_N, \hat{\mathbb{P}}_{\hat{N}})$ once, and discard all combinations $(\epsilon, \hat{\epsilon})$ with $\epsilon + \hat{\epsilon} < W(\mathbb{P}_N, \hat{\mathbb{P}}_{\hat{N}})$. The weight parameter ω of **ARO+Aux** is cross-validated from the grid $\{10^{-6}, 10^{-5}, 10^{-4}, 10^{-3}, 10^{-2}, 10^{-1}, 0, 1\}$. We fix the norms defining the feature-label metric and the adversarial attacks to ℓ_1 - and ℓ_2 -norms, respectively. The parameter κ (cf. Definition 1) is cross-validated from the grid $\{1, \sqrt{n}, n\}$, and since n is the number of features, this grid includes cases where label uncertainty is equivalent to uncertainty of a single feature ($\kappa = 1$), label uncertainty is equivalent to uncertainty of all features combined ($\kappa = n$), and an intermediary case ($\kappa = \sqrt{n}$). The case of ignoring label uncertainty ($\kappa = \infty$) is purposely not included in the grid, since one of the key reasons behind robust overfitting is that while **ARO** is equivalent to a distributionally robust model, the underlying ambiguity is only around the features (cf. our discussion in §2). All simulated adversarial attacks are worst-case ℓ_p -attacks that are instance-wise at test time, and the experiments assume we know the strength α and norm ℓ_p of the adversarial attacks.

¹In practice, distance between the unknown true and auxiliary data-generating distributions is also cross-validated in the transfer learning and domain adaptation literature (Zhong et al., 2010).

All experiments are conducted in Julia and executed on Intel Xeon 2.66GHz processors with 16GB memory in single-core mode. We use MOSEK’s exponential cone optimizer to solve all problems. To interpret the results accurately, recall that DRO+ARO and DRO+ARO+Aux are the DR models that we propose. Note also that ERM, ARO, and DRO+ARO do not utilize auxiliary data, while DRO+ARO+Aux and ARO+Aux have access to the same auxiliary datasets across all experiments (*i.e.*, we do not sample different auxiliary distributions for different methods to ensure that our comparisons are made *ceteris paribus*). Moreover, while ARO does not have access to auxiliary data, one can interpret ARO+Aux as a generalization of ARO that also has access to auxiliary datasets since it takes a mixture (with mixture weight ω) of the empirical dataset with the auxiliary dataset. For example, $\omega = 1$ would simply revise ARO by appending the empirical dataset with the auxiliary dataset.

7.1 UCI DATASETS (AUXILIARY DATA IS SYNTHETICALLY GENERATED)

We compare the out-of-sample error rates of each method on 10 UCI datasets for binary classification (Kelly et al., 2023). For each dataset, we run 10 simulations as follows: (i) Select 40% of the data as a test set ($N_{\text{te}} \propto 0.4$); (ii) Sample 25% of the remaining to form a training set ($N \propto 0.6 \cdot 0.25$); (iii) The rest ($\hat{N} \propto 0.6 \cdot 0.75$) is used to fit a synthetic generator Gaussian Copula from the SDV package (Patki et al., 2016), which is then used to generate auxiliary data. The mean errors on the test set are reported in Table 1 for ℓ_2 -attacks of strength $\alpha = 0.05$. The best error is always achieved by DRO+ARO+Aux, followed by DRO+ARO, DRO+Aux, ARO, ERM, respectively. In our appendices, we report similar results for attack strengths $\alpha \in \{0, 0.05, 0.2\}$, and share data preprocessing details and standard deviations of out-of-sample errors.

7.2 MNIST/EMNIST DATASETS (AUXILIARY DATA IS OUT-OF-DOMAIN)

We use the MNIST digits dataset (LeCun et al., 1998) to classify whether a digit is 1 or 7. For an auxiliary dataset, we use the larger EMNIST digits dataset (Cohen et al., 2017), whose authors summarize that this dataset has additional samples collected from a different group of individuals (high school students). Since EMNIST digits include MNIST digits, we remove the latter from the EMNIST dataset. We simulate the following 25 times: (i) Sample 1,000 instances from the MNIST dataset as a training set; (ii) The remaining instances in the MNIST dataset are our test set; (iii) Sample 1,000 instances from the EMNIST dataset as an auxiliary dataset. Table 2 reports the mean out-of-sample errors in various adversarial attack regimes. The results are analogous to the UCI experiments. Additionally, note that in the absence of adversarial attacks ($\alpha = 0$), DRO+ARO coincides

with the Wasserstein LR model of Shafieezadeh-Abadeh et al. (2015), and the results thus imply that even without adversarial attacks, we can improve the state-of-the-art DR model by revising its ambiguity set in light of auxiliary data.

7.3 ARTIFICIAL EXPERIMENTS (AUXILIARY DATA IS PERTURBED)

We generate empirical and auxiliary datasets by controlling their data-generating distributions (more details in the appendices). We simulate 25 cases, each with $N = 100$ training, $\hat{N} = 200$ auxiliary, and $N_{\text{te}} = 10,000$ test instances and $n = 100$ features. The performance of benchmark models with varying attacks is available in Figure 2 (left). ERM provides the worst performance, followed by ARO. The relationship between DRO+ARO and ARO+Aux is not monotonic: the former works better in larger attack regimes, conforming to the robust overfitting phenomenon. Finally, Adv+DRO+Aux always performs the best. We conduct a similar simulation for datasets with $n = 100$, and gradually increase $N = \hat{N}$ to report median ($50\% \pm 15\%$ quantiles shaded) runtimes of each method (*cf.* Figure 2, right). The fastest methods is ARO, followed by ERM, ARO+Aux, DRO+ARO, and DRO+ARO+Aux. The slowest is DRO+ARO+Aux, but the runtime scales graciously.

8 CONCLUSIONS

We formulate the distributionally robust counterpart of adversarially robust LR. Additionally, we demonstrate how to effectively utilize appropriately curated auxiliary data by intersecting Wasserstein balls. We illustrate the superiority of the proposed approach in terms of out-of-sample performance and confirm its scalability in practical settings.

From a theoretical point of view, it would be natural to extend our work to more loss functions, as is typical for DRO studies stemming from LR. To be able to optimize **Inter-ARO*** for very large-dimensional datasets, an interesting future work is to investigate first-order optimization methods that do not rely on off-the-shelf solvers. We also believe a cutting-plane method tailored for **Inter-ARO*** can also help us scale this problem for large-dimensional problems, since we would avoid monolithically optimizing a problem with $\mathcal{O}(N \cdot \hat{N})$ exponential cone constraints.

From a practical perspective, the ability to optimize **Inter-ARO*** in high-dimensional settings would also enable fine-tuning the final layer of a pre-trained neural network for binary classification, since this corresponds to logistic regression under a sigmoid activation. In our image recognition experiments, we used the MNIST dataset, as EMNIST served as a natural choice for auxiliary data. Identifying a suitable auxiliary dataset for CIFAR (Krizhevsky et al., 2009) could similarly support new experimental directions.

Table 1: Out-of-sample errors of UCI experiments with ℓ_2 -attacks of strength $\alpha = 0.05$.

Data	ERM	ARO	ARO+Aux	DRO+ARO	DRO+ARO+Aux
absent	44.02%	38.82%	35.95%	34.22%	32.64%
anneal	18.08%	16.61%	14.97%	13.50%	12.78%
audio	21.43%	21.54%	17.03%	11.76%	9.01%
breast-c	4.74%	4.93%	3.87%	3.06%	2.52%
contrac	44.14%	42.86%	40.98%	40.00%	39.65%
derma	15.97%	16.46%	13.47%	12.78%	10.84%
ecoli	16.30%	14.67%	13.26%	11.11%	9.78%
spam	11.35%	10.23%	10.16%	9.83%	9.81%
spect	33.75%	29.69%	25.78%	25.47%	21.56%
p-tumor	21.84%	20.81%	17.35%	16.18%	14.78%

Table 2: Out-of-sample errors of MNIST/EMNIST experiments with various attacks.

Attack	ERM	ARO	ARO+Aux	DRO+ARO	DRO+ARO+Aux
No attack ($\alpha = 0$)	1.55%	1.55%	1.19%	0.64%	0.53%
ℓ_1 ($\alpha = 68/255$)	2.17%	1.84%	1.33%	0.66%	0.57%
ℓ_2 ($\alpha = 128/255$)	99.93%	3.36%	2.54%	2.40%	2.12%
ℓ_∞ ($\alpha = 8/255$)	100.00%	2.60%	2.38%	2.20%	1.95%

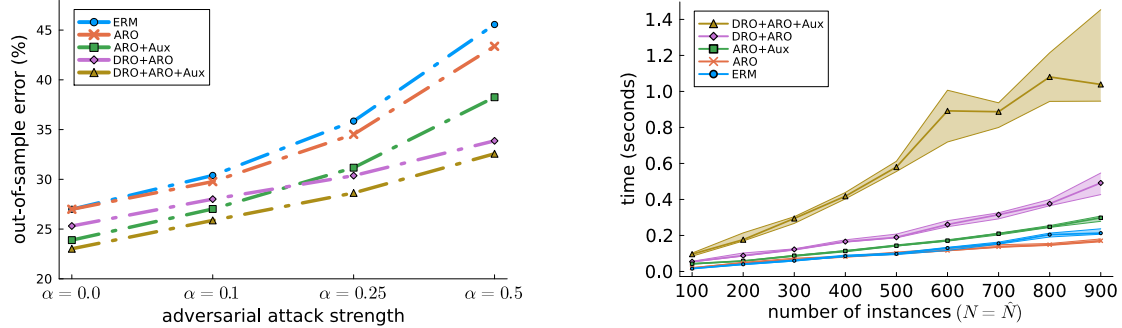


Figure 2: Out-of-sample errors under varying attack strengths (left) and runtimes under varying numbers of empirical and auxiliary instances (right) of artificial experiments.

Finally, recent breakthroughs in foundation models naturally pose the question of whether our ideas in this work apply to these models. For example, [Ye et al. \(2022\)](#) use a pre-trained language model (PLM) to generate synthetic pairs of text sequences and labels which are then used to train downstream models. It would be interesting to adapt our ideas to the text domain to explore robustness in the presence of two PLMs.

Acknowledgements

Aras Selvi’s work was done during an internship at JP Morgan AI Research. The authors gratefully acknowledge the detailed and constructive feedback provided by the anonymous reviewers and the anonymous area chair. Revising the paper has substantially improved the quality of this work.

DISCLAIMER

This paper was prepared for informational purposes by the Artificial Intelligence Research group of JPMorgan Chase & Co. and its affiliates (“JP Morgan”) and is not a product of the Research Department of JP Morgan. JP Morgan makes no representation and warranty whatsoever and disclaims all liability, for the completeness, accuracy or reliability of the information contained herein. This document is not intended as investment research or investment advice, or a recommendation, offer or solicitation for the purchase or sale of any security, financial instrument, financial product or service, or to be used in any way for evaluating the merits of participating in any transaction, and shall not constitute a solicitation under any jurisdiction or to any person, if such solicitation under such jurisdiction or to such person would be unlawful.

Bibliography

- Martin Arjovsky, Soumith Chintala, and Léon Bottou. Wasserstein generative adversarial networks. In *International Conference on Machine Learning*, 2017.
- Pranjal Awasthi, Christopher Jung, and Jamie Morgenstern. Distributionally robust data join. *arXiv:2202.05797*, 2022.
- Reza Belbasi, Aras Selvi, and Wolfram Wiesemann. It’s all in the mix: Wasserstein machine learning with mixed features. *arXiv:2312.12230*, 2023.
- Aharon Ben-Tal, Alexander Goryashko, Elana Guslitzer, and Arkadi Nemirovski. Adjustable robust solutions of uncertain linear programs. *Mathematical Programming*, 99(2):351–376, 2004.
- Aharon Ben-Tal, Laurent El Ghaoui, and Arkadi Nemirovski. *Robust Optimization*. Princeton University Press, 2009.
- Amine Bennouna and Bart Van Parys. Holistic robust data-driven decisions. *arXiv:2207.09560*, 2022.
- Amine Bennouna, Ryan Lucas, and Bart Van Parys. Certified robust neural networks: Generalization and corruption resistance. In *International Conference on Machine Learning*, 2023.
- Dimitris Bertsimas and Dick Den Hertog. *Robust and Adaptive Optimization*. Dynamic Ideas, 2022.
- Dimitris Bertsimas, Vineet Goyal, and Brian Y. Lu. A tight characterization of the performance of static solutions in two-stage adjustable robust linear optimization. *Mathematical Programming*, 150(2):281–319, 2015.
- Dimitris Bertsimas, Jack Dunn, Colin Pawlowski, and Ying Daisy Zhuo. Robust classification. *INFORMS Journal on Optimization*, 1(1):2–34, 2019.
- Battista Biggio, Iginio Corona, Davide Maiorca, Blaine Nelson, Nedim Šrđić, Pavel Laskov, Giorgio Giacinto, and Fabio Roli. Evasion attacks against machine learning at test time. In *Machine Learning and Knowledge Discovery in Databases: European Conference*, pages 387–402, 2013.
- Christopher Bishop. *Pattern Recognition and Machine Learning*. Springer, 2006.
- Tuan Anh Bui, Trung Le, Quan Tran, He Zhao, and Dinh Phung. A unified Wasserstein distributional robustness framework for adversarial training. *arXiv:2202.13437*, 2022.
- Nicholas Carlini, Anish Athalye, Nicolas Papernot, Wieland Brendel, Jonas Rauber, Dimitris Tsipras, Ian Goodfellow, Aleksander Madry, and Alexey Kurakin. On evaluating adversarial robustness. *arXiv:1902.06705*, 2019.
- Yair Carmon, Aditi Raghunathan, Ludwig Schmidt, John C. Duchi, and Percy S. Liang. Unlabeled data improves adversarial robustness. In *Advances in Neural Information Processing Systems*, 2019.
- Tianlong Chen, Zhenyu Zhang, Sijia Liu, Shiyu Chang, and Zhangyang Wang. Robust overfitting may be mitigated by properly learned smoothening. In *International Conference on Learning Representations*, 2020.
- Gregory Cohen, Saeed Afshar, Jonathan Tapson, and Andre Van Schaik. EMNIST: Extending MNIST to handwritten letters. In *International Joint Conference on Neural Networks*, 2017.
- Francesco Croce, Maksym Andriushchenko, Vikash Sehwag, Edoardo Debenedetti, Nicolas Flammarion, Mung Chiang, Prateek Mittal, and Matthias Hein. Robust-bench: a standardized adversarial robustness benchmark. *arXiv:2010.09670*, 2020.
- Erick Delage and Yinyu Ye. Distributionally robust optimization under moment uncertainty with application to data-driven problems. *Operations Research*, 58(3): 596–612, 2010.
- Victor DeMiguel and Francisco J. Nogales. Portfolio selection with robust estimation. *Operations research*, 57(3): 560–577, 2009.
- Zhun Deng, Linjun Zhang, Amirata Ghorbani, and James Zou. Improving adversarial robustness via unlabeled out-of-domain data. In *International Conference on Artificial Intelligence and Statistics*, 2021.
- Cynthia Dwork and Aaron Roth. The algorithmic foundations of differential privacy. *Foundations and Trends® in Theoretical Computer Science*, 9(3–4):211–407, 2014.
- Nicolas Fournier and Arnaud Guillin. On the rate of convergence in Wasserstein distance of the empirical measure. *Probability Theory and Related Fields*, 162(3-4):707–738, 2015.
- Natalie S. Frank and Jonathan Niles-Weed. Existence and minimax theorems for adversarial surrogate risks in binary classification. *Journal of Machine Learning Research*, 25(58):1–41, 2024.
- Rui Gao. Finite-sample guarantees for Wasserstein distributionally robust optimization: Breaking the curse of dimensionality. *Operations Research*, 71(6):2291–2306, 2023.

- Ruiqi Gao, Tianle Cai, Haochuan Li, Cho-Jui Hsieh, Liwei Wang, and Jason D. Lee. Convergence of adversarial training in overparametrized neural networks. In *Advances in Neural Information Processing Systems*, 2019.
- Clark R. Givens and Rae M. Shortt. A class of Wasserstein metrics for probability distributions. *Michigan Mathematical Journal*, 31(2):231–240, 1984.
- Ian J. Goodfellow, Jonathon Shlens, and Christian Szegedy. Explaining and harnessing adversarial examples. In *International Conference on Learning Representations*, 2015.
- Bram L. Gorissen, İhsan Yanıkoğlu, and Dick Den Hertog. A practical guide to robust optimization. *Omega*, 53: 124–137, 2015.
- Sven Gowal, Sylvestre-Alvise Rebuffi, Olivia Wiles, Florian Stimberg, Dan A. Calian, and Timothy A. Mann. Improving robustness using generated data. In *Advances in Neural Information Processing Systems*, 2021.
- Elana Guslitser. Uncertainty-immunized solutions in linear programming. Master’s thesis, Technion – Israeli Institute of Technology, 2002.
- Trevor Hastie, Robert Tibshirani, and Jerome Friedman. *The elements of statistical learning: Data mining, inference, and prediction*. Springer, 2017.
- Markelle Kelly, Rachel Longjohn, and Kolby Nottingham. The UCI machine learning repository. <https://archive.ics.uci.edu>, 2023.
- Justin Khim and Po-Ling Loh. Adversarial risk bounds via function transformation. *arXiv:1810.09519*, 2018.
- Alex Krizhevsky, Geoffrey Hinton, et al. Learning multiple layers of features from tiny images. 2009.
- Daniel Kuhn, Peyman Mohajerin Esfahani, Viet Anh Nguyen, and Soroosh Shafieezadeh-Abadeh. Wasserstein distributionally robust optimization: Theory and applications in machine learning. *INFORMS TutORials in Operations Research*, pages 130–169, 2019.
- Yann LeCun, Léon Bottou, Yoshua Bengio, and Patrick Haffner. Gradient-based learning applied to document recognition. *Proceedings of the IEEE*, 86(11):2278–2324, 1998.
- Binghui Li and Yuanzhi Li. Why clean generalization and robust overfitting both happen in adversarial training. *arXiv:2306.01271*, 2023.
- Lin Li and Michael Spratling. Understanding and combating robust overfitting via input loss landscape analysis and regularization. *Pattern Recognition*, 136:1–11, 2023.
- Runqi Lin, Chaojian Yu, and Tongliang Liu. Eliminating catastrophic overfitting via abnormal adversarial examples regularization. *Advances in Neural Information Processing Systems*, 36:67866–67885, 2023.
- Runqi Lin, Chaojian Yu, Bo Han, Hang Su, and Tongliang Liu. Layer-aware analysis of catastrophic overfitting: Revealing the pseudo-robust shortcut dependency. In Ruslan Salakhutdinov, Zico Kolter, Katherine Heller, Adrian Weller, Nuria Oliver, Jonathan Scarlett, and Felix Berkenkamp, editors, *Proceedings of the 41st International Conference on Machine Learning*, volume 235, pages 30427–30439, 21–27 Jul 2024. URL <https://proceedings.mlr.press/v235/lin24v.html>.
- Aleksander Madry, Aleksandar Makelov, Ludwig Schmidt, Dimitris Tsipras, and Adrian Vladu. Towards deep learning models resistant to adversarial attacks. In *International Conference on Learning Representations*, 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):115–166, 2018.
- MOSEK ApS. Modeling cookbook. <https://docs.mosek.com/MOSEKModelingCookbook-letter.pdf>, 2023.
- Kevin P. Murphy. *Probabilistic machine learning: an introduction*. MIT press, 2022.
- Yurii Nesterov. *Lectures on Convex Optimization*. Springer, 2018.
- Tianyu Pang, Min Lin, Xiao Yang, Jun Zhu, and Shuicheng Yan. Robustness and accuracy could be reconcilable by (proper) definition. In *International Conference on Machine Learning*, 2022.
- Neha Patki, Roy Wedge, and Kalyan Veeramachaneni. The synthetic data vault. In *IEEE International Conference on Data Science and Advanced Analytics*, 2016.
- Hoang Phan, Trung Le, Trung Phung, Anh Tuan Bui, Nhat Ho, and Dinh Phung. Global-local regularization via distributional robustness. In *International Conference on Artificial Intelligence and Statistics*, 2023.
- Muni Sreenivas Pydi and Varun Jog. The many faces of adversarial risk. In *Advances in Neural Information Processing Systems*, 2021.
- Rahul Rade and Seyed-Mohsen Moosavi-Dezfooli. Reducing excessive margin to achieve a better accuracy vs. robustness trade-off. In *International Conference on Learning Representations*, 2022.

- Aditi Raghunathan, Sang Michael Xie, Fanny Yang, John C. Duchi, and Percy Liang. Adversarial training can hurt generalization. *arXiv:1906.06032*, 2019.
- Chiara Regniet, Gauthier Gidel, and Hugo Berard. A distributional robustness perspective on adversarial training with the ∞ -Wasserstein distance. <https://openreview.net/forum?id=z7DAilcTx7>, 2022.
- Leslie Rice, Eric Wong, and Zico Kolter. Overfitting in adversarially robust deep learning. In *International Conference on Machine Learning*, 2020.
- R. Tyrrell Rockafellar. *Convex Analysis*. Princeton University Press, 1997.
- Yves Rychener, Adrián Esteban-Pérez, Juan M Morales, and Daniel Kuhn. Wasserstein distributionally robust optimization with heterogeneous data sources. *arXiv:2407.13582*, 2024.
- Vikash Sehwal, Saeed Mahloujifar, Tinashe Handina, Sihui Dai, Chong Xiang, Mung Chiang, and Prateek Mittal. Robust learning meets generative models: Can proxy distributions improve adversarial robustness? In *International Conference on Learning Representations*, 2022.
- Aras Selvi, Mohammad Reza Belbasi, Martin Haugh, and Wolfram Wiesemann. Wasserstein logistic regression with mixed features. In *Advances in Neural Information Processing Systems*, 2022.
- Ali Shafahi, Mahyar Najibi, Mohammad Amin Ghiasi, Zheng Xu, John Dickerson, Christoph Studer, Larry S Davis, Gavin Taylor, and Tom Goldstein. Adversarial training for free! In *Advances in Neural Information Processing Systems*, 2019.
- Ali Shafahi, Parsa Saadatpanah, Chen Zhu, Amin Ghiasi, Christoph Studer, David W. Jacobs, and Tom Goldstein. Adversarially robust transfer learning. In *International Conference on Learning Representations*, 2020.
- Soroosh Shafieezadeh-Abadeh, Peyman Mohajerin Esfahani, and Daniel Kuhn. Distributionally robust logistic regression. *Advances in Neural Information Processing Systems*, 2015.
- Soroosh Shafieezadeh-Abadeh, Daniel Kuhn, and Peyman Mohajerin Esfahani. Regularization via mass transportation. *Journal of Machine Learning Research*, 20(103):1–68, 2019.
- Soroosh Shafieezadeh-Abadeh, Liviu Aolaritei, Florian Dörfler, and Daniel Kuhn. New perspectives on regularization and computation in optimal transport-based distributionally robust optimization. *arXiv:2303.03900*, 2023.
- Alexander Shapiro. On duality theory of conic linear problems. *Nonconvex Optimization and its Applications*, 57: 135–155, 2001.
- Aman Sinha, Hongseok Namkoong, and John C. Duchi. Certifying some distributional robustness with principled adversarial training. In *International Conference on Learning Representations*, 2018.
- James E. Smith and Robert L. Winkler. The optimizer’s curse: Skepticism and postdecision surprise in decision analysis. *Management Science*, 52(3):311–322, 2006.
- Chuanbiao Song, Kun He, Liwei Wang, and John E. Hopcroft. Improving the generalization of adversarial training with domain adaptation. In *International Conference on Learning Representations*, 2019.
- Matthew Staib and Stefanie Jegelka. Distributionally robust deep learning as a generalization of adversarial training. In *NIPS workshop on Machine Learning and Computer Security*, 2017.
- Anirudh Subramanyam, Chrysanthos E. Gounaris, and Wolfram Wiesemann. K -adaptability in two-stage mixed-integer robust optimization. *Mathematical Programming Computation*, 12:193–224, 2020.
- Christian Szegedy, Wojciech Zaremba, Ilya Sutskever, Joan Bruna, Dumitru Erhan, Ian J. Goodfellow, and Rob Fergus. Intriguing properties of neural networks. In *International Conference on Learning Representations*, 2014.
- Neda Tanoumand, Merve Bodur, and Joe Naoum-Sawaya. Data-driven distributionally robust optimization: Intersecting ambiguity sets, performance analysis and tractability. *Optimization Online* 22567, 2023.
- Bahar Taskesen, Man-Chung Yue, Jose Blanchet, Daniel Kuhn, and Viet Anh Nguyen. Sequential domain adaptation by synthesizing distributionally robust experts. In *International Conference on Machine Learning*, 2021.
- John F. Toland. Duality in nonconvex optimization. *Journal of Mathematical Analysis and Applications*, 66(2):399–415, 1978.
- Jonathan Uesato, Brendan O’donoghue, Pushmeet Kohli, and Aaron Oord. Adversarial risk and the dangers of evaluating against weak attacks. In *International Conference on Machine Learning*, 2018.
- Jonathan Ullman and Salil Vadhan. PCPs and the hardness of generating synthetic data. *Journal of Cryptology*, 33(4):2078–2112, 2020.
- Vladimir Vapnik. *The nature of statistical learning theory*. Springer, 1999.
- Cédric Villani. *Optimal transport: Old and new*. Springer, 2009.
- Tianyu Wang, Ningyuan Chen, and Chun Wang. Contextual optimization under covariate shift: A robust approach by intersecting wasserstein balls. *arXiv:2406.02426*, 2024.

Eric Wong, Leslie Rice, and J Zico Kolter. Fast is better than free: Revisiting adversarial training. In *International Conference on Learning Representations*, 2020.

Dongxian Wu, Shu-Tao Xia, and Yisen Wang. Adversarial weight perturbation helps robust generalization. In *Advances in Neural Information Processing Systems*, 2020.

Yue Xing, Qifan Song, and Guang Cheng. Why do artificially generated data help adversarial robustness. In *Advances in Neural Information Processing Systems*, 2022a.

Yue Xing, Qifan Song, and Guang Cheng. Unlabeled data help: Minimax analysis and adversarial robustness. In *International Conference on Artificial Intelligence and Statistics*, 2022b.

İhsan Yanıkoğlu, Bram L. Gorissen, and Dick den Hertog. A survey of adjustable robust optimization. *European Journal of Operational Research*, 277(3):799–813, 2019.

Jiacheng Ye, Jiahui Gao, Quintong Li, Hang Xu, Jiangtao Feng, Zhiyong Wu, Tao Yu, and Lingpeng Kong. Zerogen: Efficient zero-shot learning via dataset generation. In *Conference on Empirical Methods in Natural Language Processing*, 2022.

Chaojian Yu, Bo Han, Li Shen, Jun Yu, Chen Gong, Mingming Gong, and Tongliang Liu. Understanding robust overfitting of adversarial training and beyond. In *International Conference on Machine Learning*, 2022.

Man-Chung Yue, Daniel Kuhn, and Wolfram Wiesemann. On linear optimization over Wasserstein balls. *Mathematical Programming*, 195(1-2):1107–1122, 2022.

Hongyang Zhang, Yaodong Yu, Jiantao Jiao, Eric Xing, Laurent El Ghaoui, and Michael Jordan. Theoretically principled trade-off between robustness and accuracy. In *International Conference on Machine Learning*, 2019.

Erheng Zhong, Wei Fan, Qiang Yang, Olivier Verscheure, and Jiangtao Ren. Cross validation framework to choose amongst models and datasets for transfer learning. In *Machine Learning and Knowledge Discovery in Databases: European Conference*, volume 6323, 2010.

Distributionally and Adversarially Robust Logistic Regression via Intersecting Wasserstein Balls

(Supplementary Material)

Aras Selvi¹ Eleonora Kreačić² Mohsen Ghassemi² Vamsi K. Potluru² Tucker Balch² Manuela Veloso²

¹Imperial Business School, Imperial College London

²JP Morgan AI Research

A NOTATION

Throughout the paper, bold lowercase letters denote vectors, while standard lowercase letters are reserved for scalars. A generic data instance is modeled as $\xi = (\mathbf{x}, y) \in \Xi := \mathbb{R}^n \times \{-1, +1\}$. For any $p > 0$, $\|\mathbf{x}\|_p$ denotes the p -norm $(\sum_{i=1}^n |x_i|^p)^{1/p}$ and $\|\mathbf{x}\|_{p^*}$ is its dual norm where $1/p + 1/p^* = 1$ with the convention of $1/1 + 1/\infty = 1$. The set of probability distributions supported on Ξ is denoted by $\mathcal{P}(\Xi)$. The Dirac measure supported on ξ is denoted by δ_ξ . The logloss is defined as $\ell_\beta(\mathbf{x}, y) = \log(1 + \exp(-y \cdot \beta^\top \mathbf{x}))$ and its associated univariate loss is $L(z) = \log(1 + \exp(-z))$ so that $L(y \cdot \beta^\top \mathbf{x}) = \ell_\beta(\mathbf{x}, y)$. The exponential cone is denoted by $\mathcal{K}_{\text{exp}} = \text{cl}(\{\omega \in \mathbb{R}^3 : \omega_1 \geq \omega_2 \cdot \exp(\omega_3/\omega_2), \omega_1 > 0, \omega_2 > 0\})$ where cl is the closure operator. The Lipschitz modulus of a univariate function f is defined as $\text{Lip}(f) := \sup_{z, z' \in \mathbb{R}} \{|f(z) - f(z')|/|z - z'| : z \neq z'\}$ whereas its effective domain is $\text{dom}(f) = \{z : f(z) < +\infty\}$. For a function $f : \mathbb{R}^n \mapsto \mathbb{R}$, its convex conjugate is $f^*(\mathbf{z}) = \sup_{\mathbf{x} \in \mathbb{R}^n} \mathbf{z}^\top \mathbf{x} - f(\mathbf{x})$. We reserve $\alpha \geq 0$ for the radii of the norms of adversarial attacks on the features and $\varepsilon \geq 0$ for the radii of distributional ambiguity sets.

B MISSING PROOFS

B.1 PROOF OF OBSERVATION 1

For any $\beta \in \mathbb{R}^n$, with standard robust optimization arguments (Ben-Tal et al., 2009; Bertsimas and Den Hertog, 2022), we can show that

$$\begin{aligned}
 & \sup_{\mathbf{z} : \|\mathbf{z}\|_p \leq \alpha} \{\ell_\beta(\mathbf{x} + \mathbf{z}, y)\} \\
 \iff & \sup_{\mathbf{z} : \|\mathbf{z}\|_p \leq \alpha} \{\log(1 + \exp(-y \cdot \beta^\top (\mathbf{x} + \mathbf{z})))\} \\
 \iff & \log \left(1 + \exp \left(\sup_{\mathbf{z} : \|\mathbf{z}\|_p \leq \alpha} \{-y \cdot \beta^\top (\mathbf{x} + \mathbf{z})\} \right) \right) \\
 \iff & \log \left(1 + \exp \left(-y \cdot \beta^\top \mathbf{x} + \alpha \cdot \sup_{\mathbf{z} : \|\mathbf{z}\|_p \leq 1} \{-y \cdot \beta^\top \mathbf{z}\} \right) \right) \\
 \iff & \log(1 + \exp(-y \cdot \beta^\top \mathbf{x} + \alpha \cdot \|\mathbf{y} \cdot \beta\|_{p^*})) \\
 \iff & \log(1 + \exp(-y \cdot \beta^\top \mathbf{x} + \alpha \cdot \|\beta\|_{p^*})),
 \end{aligned}$$

where the first step follows from the definition of logloss, the second step follows from the fact that $\log$ and $\exp$ are increasing functions, the third step takes the constant terms out of the $\sup$ problem and exploits the fact that the optimal solution of maximizing a linear function will be at an extreme point of the ℓ_p -ball, the fourth step uses the definition of dual norm, and finally, the redundant $-y \in \{-1, +1\}$ is omitted from the dual norm. We can therefore define the adversarial loss $\ell_\beta^\alpha(\mathbf{x}, y) := \log(1 + \exp(-y \cdot \beta^\top \mathbf{x} + \alpha \cdot \|\beta\|_{p^*}))$ where α models the strength of the adversary, β is the decision vector,

and $(\mathbf{x}, y)$ is an instance. Replacing $\sup_{\mathbf{z}: \|\mathbf{z}\|_p \leq \alpha} \{\ell_{\beta}(\mathbf{x} + \mathbf{z}, y)\}$ in **DR-ARO** with $\ell_{\beta}^{\alpha}(\mathbf{x}, y)$ concludes the equivalence of the optimization problem.

Furthermore, to see $\text{Lip}(L^{\alpha}) = 1$, firstly note that since $L^{\alpha}(z) = \log(1 + \exp(-z + \alpha \cdot \|\beta\|_{p^*}))$ is differentiable everywhere in z and its gradient $L^{\alpha'}$ is bounded everywhere, we have that $\text{Lip}(L^{\alpha})$ is equal to $\sup_{z \in \mathbb{R}} \{|L^{\alpha'}(z)|\}$. We thus have:

$$L^{\alpha'}(z) = \frac{-\exp(-z + \alpha \cdot \|\beta\|_{p^*})}{1 + \exp(-z + \alpha \cdot \|\beta\|_{p^*})} = \frac{-1}{1 + \exp(z - \alpha \cdot \|\beta\|_{p^*})} \in (-1, 0)$$

and $|L^{\alpha'}(z)| = [1 + \exp(z - \alpha \cdot \|\beta\|_{p^*})]^{-1} \rightarrow 1$ as $z \rightarrow -\infty$. $\square$

B.2 PROOF OF COROLLARY 1

Observation 1 lets us represent **DR-ARO** as the DR counterpart of empirical minimization of ℓ_{β}^{α} :

$$\begin{aligned} & \underset{\beta}{\text{minimize}} && \sup_{\mathbb{Q} \in \mathfrak{B}_{\varepsilon}(\mathbb{P}_N)} \mathbb{E}_{\mathbb{Q}} [\ell_{\beta}^{\alpha}(\mathbf{x}, y)] \\ & \text{subject to} && \beta \in \mathbb{R}^n. \end{aligned} \tag{4}$$

Since the univariate loss $L^{\alpha}(z) := \log(1 + \exp(-z + \alpha \cdot \|\beta\|_{p^*}))$ satisfying the identity $L^{\alpha}(\langle y \cdot \mathbf{x}, \beta \rangle) = \ell_{\beta}^{\alpha}(\mathbf{x}, y)$ is Lipschitz continuous, Theorem 14 (ii) of [Shafieezadeh-Abadeh et al. \(2019\)](#) is immediately applicable. We can therefore rewrite (4) as:

$$\begin{aligned} & \underset{\beta, \lambda, \mathbf{s}}{\text{minimize}} && \lambda \cdot \varepsilon + \frac{1}{N} \sum_{i \in [N]} s_i \\ & \text{subject to} && L^{\alpha}(\langle y^i \cdot \mathbf{x}, \beta \rangle) \leq s_i \quad \forall i \in [N] \\ & && L^{\alpha}(\langle -y^i \cdot \mathbf{x}, \beta \rangle) - \lambda \cdot \kappa \leq s_i \quad \forall i \in [N] \\ & && \text{Lip}(L^{\alpha}) \cdot \|\beta\|_{q^*} \leq \lambda \\ & && \beta \in \mathbb{R}^n, \lambda \geq 0, \mathbf{s} \in \mathbb{R}^N. \end{aligned}$$

Replacing $\text{Lip}(L^{\alpha}) = 1$ and substituting the definition of L^{α} concludes the proof. $\square$

B.3 PROOF OF PROPOSITION 1

We prove Proposition 1 by constructing the optimization problem in its statement. We will thus dualize the inner sup-problem of **Inter-ARO** for fixed β . To this end, we present a sequence of reformulations to the inner problem and then exploit strong semi-infinite duality.

By interchanging $\xi = (\mathbf{x}, y)$, we first rewrite the inner problem as

$$\begin{aligned} & \underset{\mathbb{Q}, \Pi, \widehat{\Pi}}{\text{maximize}} && \int_{\xi \in \Xi} \ell_{\beta}^{\alpha}(\xi) \mathbb{Q}(\mathrm{d}\xi) \\ & \text{subject to} && \int_{\xi, \xi' \in \Xi^2} d(\xi, \xi') \Pi(\mathrm{d}\xi, \mathrm{d}\xi') \leq \varepsilon \\ & && \int_{\xi \in \Xi} \Pi(\mathrm{d}\xi, \mathrm{d}\xi') = \mathbb{P}_N(\mathrm{d}\xi') \quad \forall \xi' \in \Xi \\ & && \int_{\xi' \in \Xi} \Pi(\mathrm{d}\xi, \mathrm{d}\xi') = \mathbb{Q}(\mathrm{d}\xi) \quad \forall \xi \in \Xi \\ & && \int_{\xi, \xi' \in \Xi^2} d(\xi, \xi') \widehat{\Pi}(\mathrm{d}\xi, \mathrm{d}\xi') \leq \widehat{\varepsilon} \\ & && \int_{\xi \in \Xi} \widehat{\Pi}(\mathrm{d}\xi, \mathrm{d}\xi') = \widehat{\mathbb{P}}_{\widehat{N}}(\mathrm{d}\xi') \quad \forall \xi' \in \Xi \\ & && \int_{\xi' \in \Xi} \widehat{\Pi}(\mathrm{d}\xi, \mathrm{d}\xi') = \mathbb{Q}(\mathrm{d}\xi) \quad \forall \xi \in \Xi \\ & && \mathbb{Q} \in \mathcal{P}(\Xi), \Pi \in \mathcal{P}(\Xi^2), \widehat{\Pi} \in \mathcal{P}(\Xi^2). \end{aligned}$$

Here, the first three constraints specify that $\mathbb{Q}$ and $\mathbb{P}_N$ have a Wasserstein distance bounded by ε from each other, modeled via their coupling Π . The latter three constraints similarly specify that $\mathbb{Q}$ and $\widehat{\mathbb{P}}_{\widehat{N}}$ are at most $\widehat{\varepsilon}$ away from each other, modeled via their coupling $\widehat{\Pi}$. As $\mathbb{Q}$ lies in the intersection of two Wasserstein balls in **Inter-ARO**, the marginal $\mathbb{Q}$ is shared between Π and $\widehat{\Pi}$. We can now substitute the third constraint into the objective and the last constraint and obtain:

$$\begin{aligned}
& \underset{\Pi, \widehat{\Pi}}{\text{maximize}} && \int_{\xi \in \Xi} \ell_{\beta}^{\alpha}(\xi) \int_{\xi' \in \Xi} \Pi(d\xi, d\xi') \\
& \text{subject to} && \int_{\xi, \xi' \in \Xi^2} d(\xi, \xi') \Pi(d\xi, d\xi') \leq \varepsilon \\
& && \int_{\xi \in \Xi} \Pi(d\xi, d\xi') = \mathbb{P}_N(d\xi') \quad \forall \xi' \in \Xi \\
& && \int_{\xi, \xi' \in \Xi^2} d(\xi, \xi') \widehat{\Pi}(d\xi, d\xi') \leq \widehat{\varepsilon} \\
& && \int_{\xi \in \Xi} \widehat{\Pi}(d\xi, d\xi') = \widehat{\mathbb{P}}_{\widehat{N}}(d\xi') \quad \forall \xi' \in \Xi \\
& && \int_{\xi' \in \Xi} \widehat{\Pi}(d\xi, d\xi') = \int_{\xi' \in \Xi} \Pi(d\xi, d\xi') \quad \forall \xi \in \Xi \\
& && \Pi \in \mathcal{P}(\Xi^2), \widehat{\Pi} \in \mathcal{P}(\Xi^2).
\end{aligned}$$

Denoting by $\mathbb{Q}^i(d\xi) := \Pi(d\xi \mid \xi^i)$ the conditional distribution of Π upon the realization of $\xi' = \xi^i$ and exploiting the fact that $\mathbb{P}_N$ is a discrete distribution supported on the N data points $\{\xi^i\}_{i \in [N]}$, we can use the marginalized representation $\Pi(d\xi, d\xi') = \frac{1}{N} \sum_{i=1}^N \delta_{\xi^i}(d\xi') \mathbb{Q}^i(d\xi)$. Similarly, we can introduce $\widehat{\mathbb{Q}}^j(d\xi) := \widehat{\Pi}(d\xi \mid \widehat{\xi}^j)$ for $\{\widehat{\xi}^j\}_{j \in [\widehat{N}]}$ to exploit the marginalized representation $\widehat{\Pi}(d\xi, d\xi') = \frac{1}{\widehat{N}} \sum_{j=1}^{\widehat{N}} \delta_{\widehat{\xi}^j}(d\xi') \widehat{\mathbb{Q}}^j(d\xi)$. By using this marginalization representation, we can use the following simplification for the objective function:

$$\int_{\xi \in \Xi} \ell_{\beta}^{\alpha}(\xi) \int_{\xi' \in \Xi} \Pi(d\xi, d\xi') = \frac{1}{N} \sum_{i=1}^N \int_{\xi \in \Xi} \ell_{\beta}^{\alpha}(\xi) \int_{\xi' \in \Xi} \delta_{\xi^i}(d\xi') \mathbb{Q}^i(d\xi) = \frac{1}{N} \sum_{i=1}^N \int_{\xi \in \Xi} \ell_{\beta}^{\alpha}(\xi) \mathbb{Q}^i(d\xi).$$

Applying analogous reformulations to the constraints leads to the following reformulation of the inner sup problem of **Inter-ARO**:

$$\begin{aligned}
& \underset{\mathbb{Q}, \widehat{\mathbb{Q}}}{\text{maximize}} && \frac{1}{N} \sum_{i=1}^N \int_{\xi \in \Xi} \ell_{\beta}^{\alpha}(\xi) \mathbb{Q}^i(d\xi) \\
& \text{subject to} && \frac{1}{N} \sum_{i=1}^N \int_{\xi \in \Xi} d(\xi, \xi^i) \mathbb{Q}^i(d\xi) \leq \varepsilon \\
& && \frac{1}{\widehat{N}} \sum_{j=1}^{\widehat{N}} \int_{\xi \in \Xi} d(\xi, \widehat{\xi}^j) \widehat{\mathbb{Q}}^j(d\xi) \leq \widehat{\varepsilon} \\
& && \frac{1}{N} \sum_{i=1}^N \mathbb{Q}^i(d\xi) = \frac{1}{\widehat{N}} \sum_{j=1}^{\widehat{N}} \widehat{\mathbb{Q}}^j(d\xi) \quad \forall \xi \in \Xi \\
& && \mathbb{Q}^i \in \mathcal{P}(\Xi), \widehat{\mathbb{Q}}^j \in \mathcal{P}(\Xi) \quad \forall i \in [N], \forall j \in [\widehat{N}].
\end{aligned}$$

We now decompose each $\mathbb{Q}^i$ into two measures corresponding to $y = \pm 1$, so that $\mathbb{Q}^i(d(\mathbf{x}, y)) = \mathbb{Q}_{+1}^i(d\mathbf{x})$ for $y = +1$ and $\mathbb{Q}^i(d(\mathbf{x}, y)) = \mathbb{Q}_{-1}^i(d\mathbf{x})$ for $y = -1$. We similarly represent each $\widehat{\mathbb{Q}}^j$ via $\widehat{\mathbb{Q}}_{+1}^j$ and $\widehat{\mathbb{Q}}_{-1}^j$ depending on y . Note that these new measures are not probability measures as they do not integrate to 1, but non-negative measures supported on $\mathbb{R}^n$

(denoted $\in \mathcal{P}_+(\mathbb{R}^n)$). We get:

$$\begin{aligned}
& \underset{\mathbb{Q}_{\pm 1}, \widehat{\mathbb{Q}}_{\pm 1}}{\text{maximize}} && \frac{1}{N} \sum_{i=1}^N \int_{\mathbf{x} \in \mathbb{R}^n} [\ell_{\boldsymbol{\beta}}^{\alpha}(\mathbf{x}, +1) \mathbb{Q}_{+1}^i(d\mathbf{x}) + \ell_{\boldsymbol{\beta}}^{\alpha}(\mathbf{x}, -1) \mathbb{Q}_{-1}^i(d\mathbf{x})] \\
& \text{subject to} && \frac{1}{N} \sum_{i=1}^N \int_{\mathbf{x} \in \mathbb{R}^n} [d((\mathbf{x}, +1), \boldsymbol{\xi}^i) \mathbb{Q}_{+1}^i(d\mathbf{x}) + d((\mathbf{x}, -1), \boldsymbol{\xi}^i) \mathbb{Q}_{-1}^i(d\mathbf{x})] \leq \varepsilon \\
& && \frac{1}{\widehat{N}} \sum_{j=1}^{\widehat{N}} \int_{\mathbf{x} \in \mathbb{R}^n} [d((\mathbf{x}, +1), \widehat{\boldsymbol{\xi}}^j) \widehat{\mathbb{Q}}_{+1}^j(d\mathbf{x}) + d((\mathbf{x}, -1), \widehat{\boldsymbol{\xi}}^j) \widehat{\mathbb{Q}}_{-1}^j(d\mathbf{x})] \leq \widehat{\varepsilon} \\
& && \int_{\mathbf{x} \in \mathbb{R}^n} \mathbb{Q}_{+1}^i(d\mathbf{x}) + \mathbb{Q}_{-1}^i(d\mathbf{x}) = 1 && \forall i \in [N] \\
& && \int_{\mathbf{x} \in \mathbb{R}^n} \widehat{\mathbb{Q}}_{+1}^j(d\mathbf{x}) + \widehat{\mathbb{Q}}_{-1}^j(d\mathbf{x}) = 1 && \forall j \in [\widehat{N}] \\
& && \frac{1}{N} \sum_{i=1}^N \mathbb{Q}_{+1}^i(d\mathbf{x}) = \frac{1}{\widehat{N}} \sum_{j=1}^{\widehat{N}} \widehat{\mathbb{Q}}_{+1}^j(d\mathbf{x}) && \forall \mathbf{x} \in \mathbb{R}^n \\
& && \frac{1}{N} \sum_{i=1}^N \mathbb{Q}_{-1}^i(d\mathbf{x}) = \frac{1}{\widehat{N}} \sum_{j=1}^{\widehat{N}} \widehat{\mathbb{Q}}_{-1}^j(d\mathbf{x}) && \forall \mathbf{x} \in \mathbb{R}^n \\
& && \mathbb{Q}_{\pm 1}^i \in \mathcal{P}_+(\mathbb{R}^n), \widehat{\mathbb{Q}}_{\pm 1}^j \in \mathcal{P}_+(\mathbb{R}^n) && \forall i \in [N], j \in [\widehat{N}].
\end{aligned}$$

Next, we explicitly write the definition of the metric $d(\cdot, \cdot)$ in the first two constraints as well as use auxiliary measures $\mathbb{A}_{\pm 1} \in \mathcal{P}_+(\mathbb{R}^n)$ to break down the last two equality constraints:

$$\begin{aligned}
& \underset{\mathbb{A}_{\pm 1}, \mathbb{Q}_{\pm 1}, \widehat{\mathbb{Q}}_{\pm 1}}{\text{maximize}} && \frac{1}{N} \sum_{i=1}^N \int_{\mathbf{x} \in \mathbb{R}^n} [\ell_{\boldsymbol{\beta}}^{\alpha}(\mathbf{x}, +1) \mathbb{Q}_{+1}^i(\mathrm{d}\mathbf{x}) + \ell_{\boldsymbol{\beta}}^{\alpha}(\mathbf{x}, -1) \mathbb{Q}_{-1}^i(\mathrm{d}\mathbf{x})] \\
& \text{subject to} && \frac{1}{N} \int_{\mathbf{x} \in \mathbb{R}^n} \left[\kappa \cdot \sum_{i \in [N]: y^i = -1} \mathbb{Q}_{+1}^i(\mathrm{d}\mathbf{x}) + \kappa \cdot \sum_{i \in [N]: y^i = +1} \mathbb{Q}_{-1}^i(\mathrm{d}\mathbf{x}) + \right. \\
& && \left. \sum_{i=1}^N \|\mathbf{x} - \mathbf{x}^i\|_q \cdot [\mathbb{Q}_{+1}^i(\mathrm{d}\mathbf{x}) + \mathbb{Q}_{-1}^i(\mathrm{d}\mathbf{x})] \right] \leq \varepsilon \\
& && \frac{1}{\widehat{N}} \int_{\mathbf{x} \in \mathbb{R}^n} \left[\kappa \cdot \sum_{j \in [N]: \widehat{y}^j = -1} \widehat{\mathbb{Q}}_{+1}^j(\mathrm{d}\mathbf{x}) + \kappa \cdot \sum_{j \in [N]: \widehat{y}^j = +1} \widehat{\mathbb{Q}}_{-1}^j(\mathrm{d}\mathbf{x}) + \right. \\
& && \left. \sum_{j=1}^{\widehat{N}} \|\mathbf{x} - \widehat{\mathbf{x}}^j\|_q \cdot [\widehat{\mathbb{Q}}_{+1}^j(\mathrm{d}\mathbf{x}) + \widehat{\mathbb{Q}}_{-1}^j(\mathrm{d}\mathbf{x})] \right] \leq \widehat{\varepsilon} \\
& && \int_{\mathbf{x} \in \mathbb{R}^n} \mathbb{Q}_{+1}^i(\mathrm{d}\mathbf{x}) + \mathbb{Q}_{-1}^i(\mathrm{d}\mathbf{x}) = 1 && \forall i \in [N] \\
& && \int_{\mathbf{x} \in \mathbb{R}^n} \widehat{\mathbb{Q}}_{+1}^j(\mathrm{d}\mathbf{x}) + \widehat{\mathbb{Q}}_{-1}^j(\mathrm{d}\mathbf{x}) = 1 && \forall j \in [\widehat{N}] \\
& && \frac{1}{N} \sum_{i=1}^N \mathbb{Q}_{+1}^i(\mathrm{d}\mathbf{x}) = \mathbb{A}_{+1}(\mathrm{d}\mathbf{x}) && \forall \mathbf{x} \in \mathbb{R}^n \\
& && \frac{1}{\widehat{N}} \sum_{j=1}^{\widehat{N}} \widehat{\mathbb{Q}}_{+1}^j(\mathrm{d}\mathbf{x}) = \mathbb{A}_{+1}(\mathrm{d}\mathbf{x}) && \forall \mathbf{x} \in \mathbb{R}^n \\
& && \frac{1}{N} \sum_{i=1}^N \mathbb{Q}_{-1}^i(\mathrm{d}\mathbf{x}) = \mathbb{A}_{-1}(\mathrm{d}\mathbf{x}) && \forall \mathbf{x} \in \mathbb{R}^n \\
& && \frac{1}{\widehat{N}} \sum_{j=1}^{\widehat{N}} \widehat{\mathbb{Q}}_{-1}^j(\mathrm{d}\mathbf{x}) = \mathbb{A}_{-1}(\mathrm{d}\mathbf{x}) && \forall \mathbf{x} \in \mathbb{R}^n \\
& && \mathbb{A}_{\pm 1} \in \mathcal{P}_+(\mathbb{R}^n), \mathbb{Q}_{\pm 1}^i \in \mathcal{P}_+(\mathbb{R}^n), \widehat{\mathbb{Q}}_{\pm 1}^j \in \mathcal{P}_+(\mathbb{R}^n) && \forall i \in [N], j \in [\widehat{N}].
\end{aligned}$$

The following semi-infinite optimization problem, obtained by standard algebraic duality, is a strong dual to the above problem since $\varepsilon, \widehat{\varepsilon} > 0$ (Shapiro, 2001).

$$\begin{aligned}
& \underset{\lambda, \hat{\lambda}, \mathbf{s}, \hat{\mathbf{s}}, p_{\pm 1}, \hat{p}_{\pm 1}}{\text{minimize}} && \frac{1}{N} \left[N\varepsilon\lambda + \hat{N}\hat{\varepsilon}\hat{\lambda} + \sum_{i=1}^N s_i + \sum_{j=1}^{\hat{N}} \hat{s}_j \right] \\
& \text{subject to} && \kappa \frac{1-y^i}{2} \lambda + \lambda \|\mathbf{x}^i - \mathbf{x}\|_q + s_i + \frac{p_{+1}(\mathbf{x})}{N} \geq \ell_{\boldsymbol{\beta}}^{\alpha}(\mathbf{x}, +1) \quad \forall i \in [N], \forall \mathbf{x} \in \mathbb{R}^n \\
& && \kappa \frac{1-\hat{y}^j}{2} \hat{\lambda} + \hat{\lambda} \|\hat{\mathbf{x}}^j - \mathbf{x}\|_q + \hat{s}_j + \frac{\hat{p}_{+1}(\mathbf{x})}{\hat{N}} \geq 0 \quad \forall j \in [\hat{N}], \forall \mathbf{x} \in \mathbb{R}^n \\
& && \kappa \frac{1+y^i}{2} \lambda + \lambda \|\mathbf{x}^i - \mathbf{x}\|_q + s_i + \frac{p_{-1}(\mathbf{x})}{N} \geq \ell_{\boldsymbol{\beta}}^{\alpha}(\mathbf{x}, -1) \quad \forall i \in [N], \forall \mathbf{x} \in \mathbb{R}^n \\
& && \kappa \frac{1+\hat{y}^j}{2} \hat{\lambda} + \hat{\lambda} \|\hat{\mathbf{x}}^j - \mathbf{x}\|_q + \hat{s}_j + \frac{\hat{p}_{-1}(\mathbf{x})}{\hat{N}} \geq 0 \quad \forall j \in [\hat{N}], \forall \mathbf{x} \in \mathbb{R}^n \\
& && p_{+1}(\mathbf{x}) + \hat{p}_{+1}(\mathbf{x}) \leq 0 \\
& && p_{-1}(\mathbf{x}) + \hat{p}_{-1}(\mathbf{x}) \leq 0 \\
& && \lambda \in \mathbb{R}_+, \hat{\lambda} \in \mathbb{R}_+, \mathbf{s} \in \mathbb{R}^N, \hat{\mathbf{s}} \in \mathbb{R}^{\hat{N}} \\
& && p_{\pm 1} : \mathbb{R}^n \mapsto \mathbb{R}, \hat{p}_{\pm 1} : \mathbb{R}^n \mapsto \mathbb{R}.
\end{aligned}$$

To eliminate the (function) variables p_{+1} and $\hat{p}_{+1}$, we first summarize the constraints they appear

$$\begin{cases}
p_{+1}(\mathbf{x}) \geq N \cdot \left[\ell_{\boldsymbol{\beta}}^{\alpha}(\mathbf{x}, +1) - s_i - \lambda \|\mathbf{x}^i - \mathbf{x}\|_q - \kappa \frac{1-y^i}{2} \lambda \right] & \forall i \in [N], \forall \mathbf{x} \in \mathbb{R}^n \\
\hat{p}_{+1}(\mathbf{x}) \geq \hat{N} \cdot \left[-\hat{s}_j - \hat{\lambda} \|\hat{\mathbf{x}}^j - \mathbf{x}\|_q - \kappa \frac{1-\hat{y}^j}{2} \hat{\lambda} \right] & \forall j \in [\hat{N}], \forall \mathbf{x} \in \mathbb{R}^n \\
p_{+1}(\mathbf{x}) + \hat{p}_{+1}(\mathbf{x}) \leq 0 & \forall \mathbf{x} \in \mathbb{R}^n,
\end{cases}$$

and notice that this system is equivalent to the epigraph-based reformulation of the following constraint

$$\begin{aligned}
& \ell_{\boldsymbol{\beta}}^{\alpha}(\mathbf{x}, +1) - s_i - \lambda \|\mathbf{x}^i - \mathbf{x}\|_q - \kappa \frac{1-y^i}{2} \lambda + \frac{\hat{N}}{N} \cdot \left[-\hat{s}_j - \hat{\lambda} \|\hat{\mathbf{x}}^j - \mathbf{x}\|_q - \kappa \frac{1-\hat{y}^j}{2} \hat{\lambda} \right] \leq 0 \\
& \forall i \in [N], \forall j \in [\hat{N}], \forall \mathbf{x} \in \mathbb{R}^n.
\end{aligned}$$

We can therefore eliminate p_{+1} and $\hat{p}_{+1}$. We can also eliminate p_{-1} and $\hat{p}_{-1}$ since we similarly have:

$$\begin{aligned}
& \begin{cases}
p_{-1}(\mathbf{x}) \geq N \cdot \left[\ell_{\boldsymbol{\beta}}^{\alpha}(\mathbf{x}, -1) - s_i - \lambda \|\mathbf{x}^i - \mathbf{x}\|_q - \kappa \frac{1+y^i}{2} \lambda \right] & \forall i \in [N], \forall \mathbf{x} \in \mathbb{R}^n \\
\hat{p}_{-1}(\mathbf{x}) \geq \hat{N} \cdot \left[-\hat{s}_j - \hat{\lambda} \|\hat{\mathbf{x}}^j - \mathbf{x}\|_q - \kappa \frac{1+\hat{y}^j}{2} \hat{\lambda} \right] & \forall j \in [\hat{N}], \forall \mathbf{x} \in \mathbb{R}^n \\
p_{-1}(\mathbf{x}) + \hat{p}_{-1}(\mathbf{x}) \leq 0 & \forall \mathbf{x} \in \mathbb{R}^n
\end{cases} \\
& \iff \ell_{\boldsymbol{\beta}}^{\alpha}(\mathbf{x}, -1) - s_i - \lambda \|\mathbf{x}^i - \mathbf{x}\|_q - \kappa \frac{1+y^i}{2} \lambda + \frac{\hat{N}}{N} \cdot \left[-\hat{s}_j - \hat{\lambda} \|\hat{\mathbf{x}}^j - \mathbf{x}\|_q - \kappa \frac{1+\hat{y}^j}{2} \hat{\lambda} \right] \leq 0 \\
& \forall i \in [N], \forall j \in [\hat{N}], \forall \mathbf{x} \in \mathbb{R}^n.
\end{aligned}$$

This trick of eliminating $p_{\pm 1}$, $\hat{p}_{\pm 1}$ is due to the auxiliary distributions $\mathbb{A}_{\pm 1}$ that we introduced; without them, the dual

problem is substantially harder to work with. We therefore obtain the following reformulation of the dual problem

$$\begin{aligned}
& \underset{\lambda, \hat{\lambda}, \mathbf{s}, \hat{\mathbf{s}}}{\text{minimize}} && \frac{1}{N} \left[N\varepsilon\lambda + \hat{N}\hat{\varepsilon}\hat{\lambda} + \sum_{i=1}^N s_i + \sum_{j=1}^{\hat{N}} \hat{s}_j \right] \\
& \text{subject to} && \sup_{\mathbf{x} \in \mathbb{R}^n} \{ \ell_{\beta}^{\alpha}(\mathbf{x}, +1) - \lambda \|\mathbf{x}^i - \mathbf{x}\|_q - \frac{\hat{N}}{N} \hat{\lambda} \|\hat{\mathbf{x}}^j - \mathbf{x}\|_q \} \leq \\
& && s_i + \kappa \frac{1 - y^i}{2} \lambda + \frac{\hat{N}}{N} \cdot \left[\hat{s}_j + \kappa \frac{1 - \hat{y}^j}{2} \hat{\lambda} \right] \quad \forall i \in [N], \forall j \in [\hat{N}] \\
& && \sup_{\mathbf{x} \in \mathbb{R}^n} \{ \ell_{\beta}^{\alpha}(\mathbf{x}, -1) - \lambda \|\mathbf{x}^i - \mathbf{x}\|_q - \frac{\hat{N}}{N} \hat{\lambda} \|\hat{\mathbf{x}}^j - \mathbf{x}\|_q \} \leq \\
& && s_i + \kappa \frac{1 + y^i}{2} \lambda + \frac{\hat{N}}{N} \cdot \left[\hat{s}_j + \kappa \frac{1 + \hat{y}^j}{2} \hat{\lambda} \right] \quad \forall i \in [N], \forall j \in [\hat{N}] \\
& && \lambda \geq 0, \hat{\lambda} \geq 0, \mathbf{s} \in \mathbb{R}_+^N, \hat{\mathbf{s}} \in \mathbb{R}_+^{\hat{N}},
\end{aligned}$$

where we replaced the $\forall \mathbf{x} \in \mathbb{R}^n$ with the worst case realizations by taking the suprema of the constraints over $\mathbf{x}$. We also added non-negativity on the definition of $\mathbf{s}$ and $\hat{\mathbf{s}}$ which is without loss of generality since this is implied by the first two constraints, which is due to the fact that in the primal reformulation the “integrates to 1” constraints (whose associated dual variables are $\mathbf{s}$ and $\hat{\mathbf{s}}$) can be written as

$$\begin{aligned}
& \int_{\mathbf{x} \in \mathbb{R}^n} \mathbb{Q}_{+1}^i(d\mathbf{x}) + \mathbb{Q}_{-1}^i(d\mathbf{x}) \leq 1 \quad \forall i \in [N], \\
& \int_{\mathbf{x} \in \mathbb{R}^n} \hat{\mathbb{Q}}_{+1}^j(d\mathbf{x}) + \hat{\mathbb{Q}}_{-1}^j(d\mathbf{x}) \leq 1 \quad \forall j \in [\hat{N}],
\end{aligned}$$

due to the objective pressure. Relabeling $\frac{\hat{N}}{N} \hat{\lambda}$ as $\hat{\lambda}$ and $\frac{\hat{N}}{N} \hat{s}_j$ as $\hat{s}_j$ simplifies the problem to:

$$\begin{aligned}
& \underset{\lambda, \hat{\lambda}, \mathbf{s}, \hat{\mathbf{s}}}{\text{minimize}} && \varepsilon\lambda + \hat{\varepsilon}\hat{\lambda} + \frac{1}{N} \sum_{i=1}^N s_i + \frac{1}{\hat{N}} \sum_{j=1}^{\hat{N}} \hat{s}_j \\
& \text{subject to} && \sup_{\mathbf{x} \in \mathbb{R}^n} \{ \ell_{\beta}^{\alpha}(\mathbf{x}, +1) - \lambda \|\mathbf{x}^i - \mathbf{x}\|_q - \hat{\lambda} \|\hat{\mathbf{x}}^j - \mathbf{x}\|_q \} \leq \\
& && s_i + \kappa \frac{1 - y^i}{2} \lambda + \hat{s}_j + \kappa \frac{1 - \hat{y}^j}{2} \hat{\lambda} \quad \forall i \in [N], \forall j \in [\hat{N}] \\
& && \sup_{\mathbf{x} \in \mathbb{R}^n} \{ \ell_{\beta}^{\alpha}(\mathbf{x}, -1) - \lambda \|\mathbf{x}^i - \mathbf{x}\|_q - \hat{\lambda} \|\hat{\mathbf{x}}^j - \mathbf{x}\|_q \} \leq \\
& && s_i + \kappa \frac{1 + y^i}{2} \lambda + \hat{s}_j + \kappa \frac{1 + \hat{y}^j}{2} \hat{\lambda} \quad \forall i \in [N], \forall j \in [\hat{N}] \\
& && \lambda \geq 0, \hat{\lambda} \geq 0, \mathbf{s} \in \mathbb{R}_+^N, \hat{\mathbf{s}} \in \mathbb{R}_+^{\hat{N}}.
\end{aligned}$$

Combining all the sup-constraints with the help of an auxiliary parameter $l \in \{-1, 1\}$ and replacing this problem with the inner problem of **Inter-ARO** concludes the proof. $\square$

B.4 PROOF OF PROPOSITION 2

We first present a technical lemma that will allow us to rewrite a specific type of difference of convex functions (DC) maximization problem that appears in the constraints of **Inter-ARO**. Rewriting such DC maximization problems is one of the key steps in reformulating Wasserstein DRO problems, and our lemma is inspired from [Shafieezadeh-Abadeh et al. \(2019, Lemma 47\)](#), [Shafieezadeh-Abadeh et al. \(2023, Theorem 3.8\)](#), and [Belbasi et al. \(2023, Lemma 1\)](#) who reformulate maximizing the difference of a convex function and a norm. Our DRO problem **Inter-ARO**, however, comprises two ambiguity sets, hence the DC term that we investigate will be the difference between a convex function and the sum of *two* norms. This requires a new analysis and we will see that **Inter-ARO** is NP-hard due to this additional difficulty.

Lemma 1. Suppose that $L : \mathbb{R} \mapsto \mathbb{R}$ is a closed convex function, and $\|\cdot\|_q$ is a norm. For vectors $\omega, \mathbf{a}, \hat{\mathbf{a}} \in \mathbb{R}^n$ and scalars $\lambda, \hat{\lambda} > 0$, we have:

$$\begin{aligned} & \sup_{\mathbf{x} \in \mathbb{R}^n} \{L(\omega^\top \mathbf{x}) - \lambda \|\mathbf{a} - \mathbf{x}\|_q - \hat{\lambda} \|\hat{\mathbf{a}} - \mathbf{x}\|_q\} \\ &= \sup_{\theta \in \text{dom}(L^*)} -L^*(\theta) + \theta \cdot \omega^\top \mathbf{a} + \theta \cdot \inf_{\mathbf{z} \in \mathbb{R}^n} \{z^\top (\hat{\mathbf{a}} - \mathbf{a}) : |\theta| \cdot \|\omega - \mathbf{z}\|_{q^*} \leq \lambda, |\theta| \cdot \|\mathbf{z}\|_{q^*} \leq \hat{\lambda}\} \end{aligned}$$

Proof. We denote by $f_\omega(\mathbf{x}) = \omega^\top \mathbf{x}$ and by g the convex function $g(\mathbf{x}) = g_1(\mathbf{x}) + g_2(\mathbf{x})$ where $g_1(\mathbf{x}) := \lambda \|\mathbf{a} - \mathbf{x}\|_q$ and $g_2(\mathbf{x}) := \hat{\lambda} \|\hat{\mathbf{a}} - \mathbf{x}\|_q$, and reformulate the sup problem as

$$\sup_{\mathbf{x} \in \mathbb{R}^n} L(\omega^\top \mathbf{x}) - g(\mathbf{x}) = \sup_{\mathbf{x} \in \mathbb{R}^n} (L \circ f_\omega)(\mathbf{x}) - g(\mathbf{x}) = \sup_{\mathbf{z} \in \mathbb{R}^n} g^*(\mathbf{z}) - (L \circ f_\omega)^*(\mathbf{z}),$$

where the first identity follows from the definition of composition and the second identity employs Toland's duality (Toland, 1978) to rewrite difference of convex functions optimization.

By using infimal convolutions (Rockafellar, 1997, Theorem 16.4), we can reformulate g^* :

$$\begin{aligned} g^*(\mathbf{z}) &= \inf_{\mathbf{z}_1, \mathbf{z}_2} \{g_1^*(\mathbf{z}_1) + g_2^*(\mathbf{z}_2) : \mathbf{z}_1 + \mathbf{z}_2 = \mathbf{z}\} \\ &= \inf_{\mathbf{z}_1, \mathbf{z}_2} \{\mathbf{z}_1^\top \mathbf{a} + \mathbf{z}_2^\top \hat{\mathbf{a}} : \mathbf{z}_1 + \mathbf{z}_2 = \mathbf{z}, \|\mathbf{z}_1\|_{q^*} \leq \lambda, \|\mathbf{z}_2\|_{q^*} \leq \hat{\lambda}\}, \end{aligned}$$

where the second step uses the definitions of $g_1^*(\mathbf{z}_1)$ and $g_2^*(\mathbf{z}_2)$. Moreover, we show

$$\begin{aligned} (L \circ f_\omega)^*(\mathbf{z}) &= \sup_{\mathbf{x} \in \mathbb{R}^n} \mathbf{z}^\top \mathbf{x} - L(\omega^\top \mathbf{x}) \\ &= \sup_{t \in \mathbb{R}, \mathbf{x} \in \mathbb{R}^n} \{\mathbf{z}^\top \mathbf{x} - L(t) : t = \omega^\top \mathbf{x}\} \\ &= \inf_{\theta \in \mathbb{R}} \sup_{t \in \mathbb{R}, \mathbf{x} \in \mathbb{R}^n} \mathbf{z}^\top \mathbf{x} - L(t) - \theta \cdot (\omega^\top \mathbf{x} - t) \\ &= \inf_{\theta \in \mathbb{R}} \sup_{t \in \mathbb{R}} \sup_{\mathbf{x} \in \mathbb{R}^n} (\mathbf{z} - \theta \cdot \omega)^\top \mathbf{x} - L(t) + \theta \cdot t \\ &= \inf_{\theta \in \mathbb{R}} \sup_{t \in \mathbb{R}} \begin{cases} -L(t) + \theta \cdot t & \text{if } \theta \cdot \omega = \mathbf{z} \\ +\infty & \text{otherwise.} \end{cases} \\ &= \inf_{\theta \in \mathbb{R}} \begin{cases} L^*(\theta) & \text{if } \theta \cdot \omega = \mathbf{z} \\ +\infty & \text{otherwise.} \end{cases} \\ &= \inf_{\theta \in \text{dom}(L^*)} \{L^*(\theta) : \theta \cdot \omega = \mathbf{z}\}, \end{aligned}$$

where the first identity follows from the definition of the convex conjugate, the second identity introduces an additional variable to make this an equality-constrained optimization problem, the third identity takes the Lagrange dual (which is a strong dual since the problem maximizes a concave objective with a single equality constraint), the fourth identity rearranges the expressions, the fifth identity exploits unboundedness of $\mathbf{x}$, the sixth identity uses the definition of convex conjugates and the final identity replaces the feasible set $\theta \in \mathbb{R}$ with the domain of L^* without loss of generality as this is an inf-problem.

Replacing the conjugates allows us to conclude that the maximization problem equals

$$\begin{aligned} & \sup_{\mathbf{z} \in \mathbb{R}^n} g^*(\mathbf{z}) + \sup_{\theta \in \text{dom}(L^*)} \{-L^*(\theta) : \theta \cdot \omega = \mathbf{z}\} \\ &= \sup_{\mathbf{z} \in \mathbb{R}^n, \theta \in \text{dom}(L^*)} \{g^*(\mathbf{z}) - L^*(\theta) : \theta \cdot \omega = \mathbf{z}\} \\ &= \sup_{\theta \in \text{dom}(L^*)} g^*(\theta \cdot \omega) - L^*(\theta) \\ &= \sup_{\theta \in \text{dom}(L^*)} -L^*(\theta) + \inf_{\mathbf{z}_1, \mathbf{z}_2 \in \mathbb{R}^n} \{\mathbf{z}_1^\top \mathbf{a} + \mathbf{z}_2^\top \hat{\mathbf{a}} : \mathbf{z}_1 + \mathbf{z}_2 = \theta \cdot \omega, \|\mathbf{z}_1\|_{q^*} \leq \lambda, \|\mathbf{z}_2\|_{q^*} \leq \hat{\lambda}\} \\ &= \sup_{\theta \in \text{dom}(L^*)} -L^*(\theta) + \theta \cdot \inf_{\mathbf{z}_1, \mathbf{z}_2 \in \mathbb{R}^n} \{\mathbf{z}_1^\top \mathbf{a} + \mathbf{z}_2^\top \hat{\mathbf{a}} : \mathbf{z}_1 + \mathbf{z}_2 = \omega, |\theta| \cdot \|\mathbf{z}_1\|_{q^*} \leq \lambda, |\theta| \cdot \|\mathbf{z}_2\|_{q^*} \leq \hat{\lambda}\} \\ &= \sup_{\theta \in \text{dom}(L^*)} -L^*(\theta) + \theta \cdot \omega^\top \mathbf{a} + \theta \cdot \inf_{\mathbf{z} \in \mathbb{R}^n} \{z^\top (\hat{\mathbf{a}} - \mathbf{a}) : |\theta| \cdot \|\omega - \mathbf{z}\|_{q^*} \leq \lambda, |\theta| \cdot \|\mathbf{z}\|_{q^*} \leq \hat{\lambda}\}. \end{aligned}$$

Here, the first identity follows from writing the problem as a single maximization problem, the second identity follows from the equality constraint, the third identity follows from the definition of the conjugate g^* , the fourth identity is due to relabeling $z_1 = \theta \cdot z_1$ and $z_2 = \theta \cdot z_2$, and the fifth identity is due to a variable change ($z_1 = \omega - z_2$ relabeled as z). $\square$

DC maximization terms similar to the one dealt by Lemma 1 appear on the left-hand side of the constraints of **Inter-ARO** (cf. formulation in Proposition 1). These constraints would admit a tractable reformulation for the case without auxiliary data because the $\inf$ -term in the reformulation presented in Lemma 1 does not appear in such cases. To see this, eliminate the second norm (the one associated with auxiliary data) by taking $\hat{\lambda} = 0$, which will cause the constraint $|\theta| \cdot \|z\|_{q^*} \leq \hat{\lambda}$ to force $z = 0$, and the alternative formulation will thus be:

$$\begin{aligned} & \begin{cases} \sup_{\theta \in \text{dom}(L^*)} \{-L^*(\theta) + \theta \cdot \omega^\top \mathbf{a}\} & \text{if } \sup_{\theta \in \text{dom}(L^*)} \{|\theta|\} \cdot \|z\|_{q^*} \leq \lambda \\ +\infty & \text{otherwise} \end{cases} \\ = & \begin{cases} L(\omega^\top \mathbf{a}) & \text{if } \text{Lip}(L) \cdot \|z\|_{q^*} \leq \lambda \\ +\infty & \text{otherwise,} \end{cases} \end{aligned}$$

where we used the fact that $L = L^{**}$ and $\sup_{\theta \in \text{dom}(L)} |\theta| = \text{Lip}(L)$ since L is closed convex (Rockafellar, 1997, Corollary 13.3.3). Hence, the DC maximization can be represented with a convex function with an additional convex inequality, making the constraints tractable for the case without auxiliary data. For the case with auxiliary data, however, the $\sup_\theta \inf_z$ structure makes these constraints equivalent to two-stage robust constraints (with uncertain parameter θ and adjustable variable z), bringing an adjustable robust optimization (Ben-Tal et al., 2004; Yanıkoğlu et al., 2019) perspective to **Inter-ARO**. By using the univariate representation $\ell_\beta^\alpha(x, y) = L^\alpha(y \cdot \beta^\top x)$, **Inter-ARO** can be written as

$$\begin{aligned} \underset{\beta, \lambda, \hat{\lambda}, \mathbf{s}, \hat{\mathbf{s}}}{\text{minimize}} \quad & \varepsilon \lambda + \hat{\varepsilon} \hat{\lambda} + \frac{1}{N} \sum_{j=1}^N s_j + \frac{1}{\hat{N}} \sum_{i=1}^{\hat{N}} \hat{s}_i \\ \text{subject to} \quad & \sup_{\mathbf{x} \in \mathbb{R}^n} \{L^\alpha(\beta^\top \mathbf{x}) - \lambda \|\mathbf{x}^i - \mathbf{x}\|_q - \hat{\lambda} \|\hat{\mathbf{x}}^j - \mathbf{x}\|_q\} \leq \\ & s_i + \kappa \frac{1 - y^i}{2} \lambda + \hat{s}_j + \kappa \frac{1 - \hat{y}^j}{2} \hat{\lambda} \quad \forall i \in [N], \forall j \in [\hat{N}] \\ & \sup_{\mathbf{x} \in \mathbb{R}^n} \{L^\alpha(-\beta^\top \mathbf{x}) - \lambda \|\mathbf{x}^i - \mathbf{x}\|_q - \hat{\lambda} \|\hat{\mathbf{x}}^j - \mathbf{x}\|_q\} \leq \\ & s_i + \kappa \frac{1 + y^i}{2} \lambda + \hat{s}_j + \kappa \frac{1 + \hat{y}^j}{2} \hat{\lambda} \quad \forall i \in [N], \forall j \in [\hat{N}] \\ & \beta \in \mathbb{R}^n, \lambda \geq 0, \hat{\lambda} \geq 0, \mathbf{s} \in \mathbb{R}_+^N, \hat{\mathbf{s}} \in \mathbb{R}_+^{\hat{N}}, \end{aligned}$$

and applying Lemma 1 to the left-hand side of the constraints gives:

$$\begin{aligned} \underset{\beta, \lambda, \hat{\lambda}, \mathbf{s}, \hat{\mathbf{s}}}{\text{minimize}} \quad & \varepsilon \lambda + \hat{\varepsilon} \hat{\lambda} + \frac{1}{N} \sum_{j=1}^N s_j + \frac{1}{\hat{N}} \sum_{i=1}^{\hat{N}} \hat{s}_i \\ \text{subject to} \quad & \sup_{\theta \in \text{dom}(L^*)} -L^{\alpha*}(\theta) + \theta \cdot \beta^\top \mathbf{x}^i + \theta \cdot \inf_{\mathbf{z} \in \mathbb{R}^n} \{\mathbf{z}^\top (\hat{\mathbf{x}}^j - \mathbf{x}^i) : |\theta| \cdot \|\beta - \mathbf{z}\|_{q^*} \leq \lambda, |\theta| \cdot \|\mathbf{z}\|_{q^*} \leq \hat{\lambda}\} \leq \\ & s_i + \kappa \frac{1 - y^i}{2} \lambda + \hat{s}_j + \kappa \frac{1 - \hat{y}^j}{2} \hat{\lambda} \quad \forall i \in [N], \forall j \in [\hat{N}] \\ & \sup_{\theta \in \text{dom}(L^*)} -L^{\alpha*}(\theta) - \theta \cdot \beta^\top \mathbf{x}^i + \theta \cdot \inf_{\mathbf{z} \in \mathbb{R}^n} \{\mathbf{z}^\top (\hat{\mathbf{x}}^j - \mathbf{x}^i) : |\theta| \cdot \|-\beta - \mathbf{z}\|_{q^*} \leq \lambda, |\theta| \cdot \|\mathbf{z}\|_{q^*} \leq \hat{\lambda}\} \leq \\ & s_i + \kappa \frac{1 + y^i}{2} \lambda + \hat{s}_j + \kappa \frac{1 + \hat{y}^j}{2} \hat{\lambda} \quad \forall i \in [N], \forall j \in [\hat{N}] \\ & \beta \in \mathbb{R}^n, \lambda \geq 0, \hat{\lambda} \geq 0, \mathbf{s} \in \mathbb{R}_+^N, \hat{\mathbf{s}} \in \mathbb{R}_+^{\hat{N}}. \end{aligned} \tag{5}$$

Which, equivalently, can be written as the following problem with $2N \cdot \hat{N}$ two-stage robust constraints:

$$\begin{aligned}
& \underset{\beta, \lambda, \hat{\lambda}, s, \hat{s}}{\text{minimize}} && \varepsilon\lambda + \hat{\varepsilon}\hat{\lambda} + \frac{1}{N} \sum_{j=1}^N s_j + \frac{1}{\hat{N}} \sum_{i=1}^{\hat{N}} \hat{s}_i \\
& \text{subject to} && \left[\forall \theta \in \text{dom}(L^*), \exists \mathbf{z} \in \mathbb{R}^n : \begin{cases} -L^{\alpha*}(\theta) + \theta \cdot \beta^\top \mathbf{x}^i + \theta \cdot \mathbf{z}^\top (\hat{\mathbf{x}}^j - \mathbf{x}^i) \leq s_i + \kappa \frac{1-y^i}{2} \lambda + \hat{s}_j + \kappa \frac{1-\hat{y}^j}{2} \hat{\lambda} \\ |\theta| \cdot \|\beta - \mathbf{z}\|_{q^*} \leq \lambda \\ |\theta| \cdot \|\mathbf{z}\|_{q^*} \leq \hat{\lambda} \end{cases} \right] \\
& && \forall i \in [N], \forall j \in [\hat{N}] \\
& && \left[\forall \theta \in \text{dom}(L^*), \exists \mathbf{z} \in \mathbb{R}^n : \begin{cases} -L^{\alpha*}(\theta) - \theta \cdot \beta^\top \mathbf{x}^i + \theta \cdot \mathbf{z}^\top (\hat{\mathbf{x}}^j - \mathbf{x}^i) \leq s_i + \kappa \frac{1+y^i}{2} \lambda + \hat{s}_j + \kappa \frac{1+\hat{y}^j}{2} \hat{\lambda} \\ |\theta| \cdot \|-\beta - \mathbf{z}\|_{q^*} \leq \lambda \\ |\theta| \cdot \|\mathbf{z}\|_{q^*} \leq \hat{\lambda} \end{cases} \right] \\
& && \forall i \in [N], \forall j \in [\hat{N}]
\end{aligned}$$

$$\beta \in \mathbb{R}^n, \lambda \geq 0, \hat{\lambda} \geq 0, \mathbf{s} \in \mathbb{R}_+^N, \hat{\mathbf{s}} \in \mathbb{R}_+^{\hat{N}}. \quad (\text{Inter-adjustable})$$

By using adjustable robust optimization theory, we show that this problem is NP-hard even in the simplest setting. To this end, take $N = \hat{N} = 1$ as well as $\kappa = 0$; the formulation presented in Proposition 1 reduces to:

$$\begin{aligned}
& \underset{\beta, \lambda, \hat{\lambda}, s, \hat{s}}{\text{minimize}} && \varepsilon\lambda + \hat{\varepsilon}\hat{\lambda} + s + \hat{s} \\
& \text{subject to} && \sup_{\mathbf{x} \in \mathbb{R}^n} \{ \ell_\beta^\alpha(\mathbf{x}, l) - \lambda \|\mathbf{x}^1 - \mathbf{x}\|_q - \hat{\lambda} \|\hat{\mathbf{x}}^1 - \mathbf{x}\|_q \} \leq s_1 + \hat{s}_1 \quad \forall l \in \{-1, 1\} \\
& && \beta \in \mathbb{R}^n, \lambda \geq 0, \hat{\lambda} \geq 0, s \geq 0, \hat{s} \geq 0.
\end{aligned}$$

The worst case realization of $l \in \{-1, 1\}$ will always make $\ell_\beta^\alpha(\mathbf{x}, l) = \log(1 + \exp(-l \cdot \beta^\top \mathbf{x} + \alpha \cdot \|\beta\|_{p^*}))$ equal to $\varsigma_\beta^\alpha(\mathbf{x}) = \log(1 + \exp(|l \cdot \beta^\top \mathbf{x}| + \alpha \cdot \|\beta\|_{p^*}))$, where ς inherits similar properties from ℓ : it is convex in β and its univariate representation S^α has the same Lipschitz constant with L^α . We can thus represent the above problem as

$$\begin{aligned}
& \underset{\beta, \lambda, \hat{\lambda}, s, \hat{s}}{\text{minimize}} && \varepsilon\lambda + \hat{\varepsilon}\hat{\lambda} + s + \hat{s} \\
& \text{subject to} && \sup_{\mathbf{x} \in \mathbb{R}^n} \{ S^\alpha(\beta^\top \mathbf{x}) - \lambda \|\mathbf{x}^1 - \mathbf{x}\|_q - \hat{\lambda} \|\hat{\mathbf{x}}^1 - \mathbf{x}\|_q \} \leq s + \hat{s} \\
& && \beta \in \mathbb{R}^n, \lambda \geq 0, \hat{\lambda} \geq 0, s \geq 0, \hat{s} \geq 0.
\end{aligned}$$

Substituting $s + \hat{s}$ into the objective (due to the objective pressure) allows us to reformulate the above problem as

$$\begin{aligned}
& \underset{\beta, \lambda, \hat{\lambda}}{\text{minimize}} && \varepsilon\lambda + \hat{\varepsilon}\hat{\lambda} + \sup_{\mathbf{x} \in \mathbb{R}^n} \{ S^\alpha(\beta^\top \mathbf{x}) - \lambda \|\mathbf{x}^1 - \mathbf{x}\|_q - \hat{\lambda} \|\hat{\mathbf{x}}^1 - \mathbf{x}\|_q \} \\
& \text{subject to} && \beta \in \mathbb{R}^n, \lambda \geq 0, \hat{\lambda} \geq 0,
\end{aligned} \quad (6)$$

and an application of Lemma 1 leads us to the following reformulation:

$$\inf_{\substack{\beta \in \mathbb{R}^n \\ \lambda \geq 0, \hat{\lambda} \geq 0}} \sup_{\theta \in \text{dom}(S^*)} \inf_{\mathbf{z} \in \mathbb{R}^n} \left\{ \varepsilon\lambda + \hat{\varepsilon}\hat{\lambda} - S^{\alpha*}(\theta) + \underbrace{\theta \cdot \beta^\top \mathbf{x}^1 + \theta \cdot \mathbf{z}^\top (\hat{\mathbf{x}}^1 - \mathbf{x}^1)}_{(1)} : \underbrace{|\theta| \cdot \|\beta - \mathbf{z}\|_{q^*} \leq \lambda, |\theta| \cdot \|\mathbf{z}\|_{q^*} \leq \hat{\lambda}}_{(2)} \right\}.$$

The objective term (1) has a product of the uncertain parameter θ and the adjustable variable $\mathbf{z}$, and even when (2) is linear such as in the case of $q = 1$ the product of the uncertain parameter with both the decision variable β and the adjustable variable $\mathbf{z}$ still appear since:

$$|\theta| \cdot \|\beta - \mathbf{z}\|_\infty \leq \lambda \iff -\lambda \leq \theta\beta - \theta\mathbf{z} \leq \lambda.$$

This reduces problem (6) to a generic two-stage robust optimization problem with random recourse (Subramanyam et al., 2020, Problem 1) which is proven to be NP-hard even if $S^{\alpha*}$ was constant (Guslitser, 2002). $\square$

B.5 PROOF OF THEOREM 1

Consider the reformulation **Inter-adjustable** of **Inter-ARO** that we introduced in the proof of Proposition 2. For any $i \in [N]$ and $j \in [\hat{N}]$, the corresponding constraint in the first group of ‘adjustable robust’ ($\forall, \exists$) constraints will be:

$$\forall \theta \in \text{dom}(L^*), \exists \mathbf{z} \in \mathbb{R}^n : \begin{cases} -L^{\alpha*}(\theta) + \theta \cdot \boldsymbol{\beta}^\top \mathbf{x}^i + \theta \cdot \mathbf{z}^\top (\hat{\mathbf{x}}^j - \mathbf{x}^i) \leq s_i + \kappa \frac{1-y^i}{2} \lambda + \hat{s}_j + \kappa \frac{1-\hat{y}^j}{2} \hat{\lambda} \\ |\theta| \cdot \|\boldsymbol{\beta} - \mathbf{z}\|_{q^*} \leq \lambda \\ |\theta| \cdot \|\mathbf{z}\|_{q^*} \leq \hat{\lambda}. \end{cases}$$

By changing the order of $\forall$ and $\exists$, we obtain:

$$\exists \mathbf{z} \in \mathbb{R}^n, \forall \theta \in \text{dom}(L^*) : \begin{cases} -L^{\alpha*}(\theta) + \theta \cdot \boldsymbol{\beta}^\top \mathbf{x}^i + \theta \cdot \mathbf{z}^\top (\hat{\mathbf{x}}^j - \mathbf{x}^i) \leq s_i + \kappa \frac{1-y^i}{2} \lambda + \hat{s}_j + \kappa \frac{1-\hat{y}^j}{2} \hat{\lambda} \\ |\theta| \cdot \|\boldsymbol{\beta} - \mathbf{z}\|_{q^*} \leq \lambda \\ |\theta| \cdot \|\mathbf{z}\|_{q^*} \leq \hat{\lambda}. \end{cases}$$

Notice that this is a safe approximation, since any fixed $\mathbf{z}$ satisfying the latter system is a feasible static solution in the former system, meaning that for every realization of θ in the first system, the inner $\exists \mathbf{z}$ can always ‘play’ the same $\mathbf{z}$ that is feasible in the latter system (hence the latter is named the *static* relaxation, [Bertsimas et al. 2015](#)). In the relaxed system, we can drop $\forall \theta$ and keep its worst-case realization instead:

$$\exists \mathbf{z} \in \mathbb{R}^n : \begin{cases} \sup_{\theta \in \text{dom}(L^*)} \{-L^{\alpha*}(\theta) + \theta \cdot \boldsymbol{\beta}^\top \mathbf{x}^i + \theta \cdot \mathbf{z}^\top (\hat{\mathbf{x}}^j - \mathbf{x}^i)\} \leq s_i + \kappa \frac{1-y^i}{2} \lambda + \hat{s}_j + \kappa \frac{1-\hat{y}^j}{2} \hat{\lambda} \\ \sup_{\theta \in \text{dom}(L^*)} \{|\theta|\} \cdot \|\boldsymbol{\beta} - \mathbf{z}\|_{q^*} \leq \lambda \\ \sup_{\theta \in \text{dom}(L^*)} \{|\theta|\} \cdot \|\mathbf{z}\|_{q^*} \leq \hat{\lambda}. \end{cases}$$

The term $\sup_{\theta \in \text{dom}(L^*)} \{-L^{\alpha*}(\theta) + \theta \cdot \boldsymbol{\beta}^\top \mathbf{x}^i + \theta \cdot \mathbf{z}^\top (\hat{\mathbf{x}}^j - \mathbf{x}^i)\}$ is the definition of the biconjugate $L^{\alpha**}(\boldsymbol{\beta}^\top \mathbf{x}^i + \mathbf{z}^\top (\hat{\mathbf{x}}^j - \mathbf{x}^i))$. Since L^α is a closed convex function, we have $L^{\alpha**} = L^\alpha$ ([Rockafellar, 1997](#), Corollary 12.2.1). Moreover, $\sup_{\theta \in \text{dom}(L^*)} \{|\theta|\}$ is an alternative representation of the Lipschitz constant of the function L^α ([Rockafellar, 1997](#), Corollary 13.3.3), which is equal to 1 as we showed earlier. The adjustable robust constraint thus reduces to:

$$\exists \mathbf{z} \in \mathbb{R}^n : \begin{cases} L^\alpha(\boldsymbol{\beta}^\top \mathbf{x}^i + \mathbf{z}^\top (\hat{\mathbf{x}}^j - \mathbf{x}^i)) \leq s_i + \kappa \frac{1-y^i}{2} \lambda + \hat{s}_j + \kappa \frac{1-\hat{y}^j}{2} \hat{\lambda} \\ \|\boldsymbol{\beta} - \mathbf{z}\|_{q^*} \leq \lambda \\ \|\mathbf{z}\|_{q^*} \leq \hat{\lambda} \end{cases}$$

as a result of the static relaxation. This relaxed reformulation applies to all $i \in [N]$ and $j \in [\hat{N}]$ as well as to the second group of adjustable robust constraints analogously. Replacing each constraint of **Inter-adjustable** with this system concludes the proof. $\square$

B.6 PROOF OF COROLLARY 2

To prove the first statement, take $\hat{\lambda} = 0$ and observe the constraint $\|\mathbf{z}_{ij}^l\|_{q^*} \leq \hat{\lambda}$ implies $\mathbf{z}_{ij}^l = \mathbf{0}$ for all $l \in \{-1, 1\}$, $i \in [N]$, $j \in [\hat{N}]$. The optimization problem can thus be written without those variables:

$$\begin{aligned} & \underset{\boldsymbol{\beta}, \lambda, \mathbf{s}, \hat{\mathbf{s}}}{\text{minimize}} && \varepsilon \lambda + \frac{1}{N} \sum_{i=1}^N s_i + \frac{1}{\hat{N}} \sum_{j=1}^{\hat{N}} \hat{s}_j \\ & \text{subject to} && L^\alpha(l \boldsymbol{\beta}^\top \mathbf{x}^i) \leq s_i + \kappa \frac{1-ly^i}{2} \lambda + \hat{s}_j \quad \forall l \in \{-1, 1\}, \forall i \in [N], \forall j \in [\hat{N}] \\ & && \|\boldsymbol{\beta}\|_{q^*} \leq \lambda \\ & && \boldsymbol{\beta} \in \mathbb{R}^n, \lambda \geq 0, \mathbf{s} \in \mathbb{R}_+^N, \hat{\mathbf{s}} \in \mathbb{R}_+^{\hat{N}}. \end{aligned}$$

Notice that optimal solutions should satisfy $\hat{s}_j = \hat{s}_{j'}$ for all $j, j' \in [N]$. To see this, assume for contradiction that $\exists j, j' \in [N]$ such that $\hat{s}_j < \hat{s}_{j'}$. If a constraint indexed with (l, i, j) for arbitrary $l \in \{-1, 1\}$ and $i \in [N]$ is feasible, it means the

constraint indexed with (l, i, j') cannot be tight given that these constraints are identical except for the $\hat{s}_j$ or $\hat{s}_{j'}$ appearing on the right hand side. Hence, such a solution cannot be optimal as this is a minimization problem, and updating $\hat{s}_{j'}$ as $\hat{s}_j$ preserves the feasibility of the problem while decreasing the objective value. We can thus use a single variable $\tau \in \mathbb{R}_+$ and rewrite the problem as

$$\begin{aligned} & \underset{\boldsymbol{\beta}, \lambda, \mathbf{s}, \hat{\mathbf{s}}}{\text{minimize}} && \varepsilon \lambda + \frac{1}{N} \sum_{i=1}^N (s_i + \tau) \\ & \text{subject to} && L^\alpha(\boldsymbol{\beta}^\top \mathbf{x}^i) \leq s_i + \kappa \frac{1-y^i}{2} \lambda + \tau \quad \forall i \in [N] \\ & && L^\alpha(-\boldsymbol{\beta}^\top \mathbf{x}^i) \leq s_i + \kappa \frac{1+y^i}{2} \lambda + \tau \quad \forall i \in [N] \\ & && \|\boldsymbol{\beta}\|_{q^*} \leq \lambda \\ & && \boldsymbol{\beta} \in \mathbb{R}^n, \lambda \geq 0, \mathbf{s} \in \mathbb{R}_+^N, \hat{\mathbf{s}} \in \mathbb{R}_+^{\hat{N}}, \end{aligned}$$

where we also eliminated the index $l \in \{-1, 1\}$ by writing the constraints explicitly. Since s_i and τ both appear as $s_i + \tau$ in this problem, we can use a variable change where we relabel $s_i + \tau$ as s_i (or, equivalently set $\tau = 0$ without any optimality loss). Moreover, the constraints with index $i \in [N]$ are

$$\begin{cases} L^\alpha(\boldsymbol{\beta}^\top \mathbf{x}^i) \leq s_i + \tau \\ L^\alpha(-\boldsymbol{\beta}^\top \mathbf{x}^i) \leq s_i + \kappa \lambda + \tau \end{cases} = \begin{cases} L^\alpha(y^i \cdot \boldsymbol{\beta}^\top \mathbf{x}^i) \leq s_i + \tau \\ L^\alpha(-y^i \cdot \boldsymbol{\beta}^\top \mathbf{x}^i) \leq s_i + \kappa \lambda + \tau \end{cases}$$

if $y^i = 1$, and similarly they are

$$\begin{cases} L^\alpha(\boldsymbol{\beta}^\top \mathbf{x}^i) \leq s_i + \kappa \lambda + \tau \\ L^\alpha(-\boldsymbol{\beta}^\top \mathbf{x}^i) \leq s_i + \tau \end{cases} = \begin{cases} L^\alpha(-y^i \cdot \boldsymbol{\beta}^\top \mathbf{x}^i) \leq s_i + \kappa \lambda + \tau \\ L^\alpha(y^i \cdot \boldsymbol{\beta}^\top \mathbf{x}^i) \leq s_i + \tau \end{cases}$$

if $y^i = -1$. Since these are identical, the problem can finally be written as

$$\begin{aligned} & \underset{\boldsymbol{\beta}, \lambda, \mathbf{s}}{\text{minimize}} && \varepsilon \lambda + \frac{1}{N} \sum_{i=1}^N s_i \\ & \text{subject to} && \log(1 + \exp(-y^i \cdot \boldsymbol{\beta}^\top \mathbf{x}^i + \alpha \cdot \|\boldsymbol{\beta}\|_{p^*})) \leq s_i \quad \forall i \in [N] \\ & && \log(1 + \exp(y^i \cdot \boldsymbol{\beta}^\top \mathbf{x}^i + \alpha \cdot \|\boldsymbol{\beta}\|_{p^*})) - \lambda \kappa \leq s_i \quad \forall i \in [N] \\ & && \|\boldsymbol{\beta}\|_{q^*} \leq \lambda \\ & && \boldsymbol{\beta} \in \mathbb{R}^n, \lambda \geq 0, \mathbf{s} \in \mathbb{R}_+^N, \end{aligned}$$

where we also used the definition of L^α . This problem is identical to **DR-ARO**, which means that feasible solutions of **DR-ARO** are feasible for **Inter-ARO*** if the additional variables $(\hat{\lambda}, \hat{\mathbf{s}}, \mathbf{z}_{ij}^l)$ are set to zero, concluding the first statement of the corollary.

The second statement is immediate since $\hat{\varepsilon} \rightarrow \infty$ forces $\hat{\lambda} = 0$ due to the term $\hat{\varepsilon} \hat{\lambda}$ in the objective of **Inter-ARO***, and this proof shows in such a case **Inter-ARO*** reduces to **DR-ARO** (which is identical to **Inter-ARO** when $\varepsilon \rightarrow \infty$ by definition). $\square$

B.7 PROOF OF OBSERVATION 2

By standard linearity arguments and from the definition of $\mathbb{Q}_{\text{mix}}$, we have

$$\begin{aligned}
& \mathbb{E}_{\mathbb{Q}_{\text{mix}}} \left[\sup_{\mathbf{z} \in \mathcal{B}_p(\alpha)} \{ \ell_{\beta}(\mathbf{x} + \mathbf{z}, y) \} \right] \\
& \iff \int_{(\mathbf{x}, y) \in \mathbb{R}^n \times \{-1, +1\}} \sup_{\mathbf{z} \in \mathcal{B}_p(\alpha)} \{ \ell_{\beta}(\mathbf{x} + \mathbf{z}, y) \} d\mathbb{Q}_{\text{mix}}((\mathbf{x}, y)) \\
& \iff \frac{N}{N + w\widehat{N}} \int_{(\mathbf{x}, y) \in \mathbb{R}^n \times \{-1, +1\}} \sup_{\mathbf{z} \in \mathcal{B}_p(\alpha)} \{ \ell_{\beta}(\mathbf{x} + \mathbf{z}, y) \} d\mathbb{P}_N((\mathbf{x}, y)) + \\
& \quad \frac{w\widehat{N}}{N + w\widehat{N}} \int_{(\mathbf{x}, y) \in \mathbb{R}^n \times \{-1, +1\}} \sup_{\mathbf{z} \in \mathcal{B}_p(\alpha)} \{ \ell_{\beta}(\mathbf{x} + \mathbf{z}, y) \} d\widehat{\mathbb{P}}_{\widehat{N}}((\mathbf{x}, y)) \\
& \iff \frac{N}{N + w\widehat{N}} \cdot \frac{1}{N} \sum_{i \in [N]} \sup_{\mathbf{z}^i \in \mathcal{B}_p(\alpha)} \{ \ell_{\beta}(\mathbf{x}^i + \mathbf{z}^i, y^i) \} + \frac{w\widehat{N}}{N + w\widehat{N}} \cdot \frac{1}{\widehat{N}} \sum_{j \in [\widehat{N}]} \sup_{\mathbf{z}^j \in \mathcal{B}_p(\alpha)} \{ \ell_{\beta}(\widehat{\mathbf{x}}^j + \mathbf{z}^j, \widehat{y}^j) \} \\
& \iff \frac{1}{N + w\widehat{N}} \left[\sum_{i \in [N]} \sup_{\mathbf{z}^i \in \mathcal{B}_p(\alpha)} \{ \ell_{\beta}(\mathbf{x}^i + \mathbf{z}^i, y^i) \} + w \cdot \sum_{j \in [\widehat{N}]} \sup_{\mathbf{z}^j \in \mathcal{B}_p(\alpha)} \{ \ell_{\beta}(\widehat{\mathbf{x}}^j + \mathbf{z}^j, \widehat{y}^j) \} \right],
\end{aligned}$$

which coincides with the objective function of (2). Since we have

$$\mathbb{E}_{\mathbb{Q}_{\text{mix}}} \left[\sup_{\mathbf{z} \in \mathcal{B}_p(\alpha)} \{ \ell_{\beta}(\mathbf{x} + \mathbf{z}, y) \} \right] = \mathbb{E}_{\mathbb{Q}_{\text{mix}}} [\ell_{\beta}^{\alpha}(\mathbf{x}, y)]$$

we can conclude the proof. $\square$

B.8 PROOF OF PROPOSITION 3

We first prove auxiliary results on mixture distributions. To this end, denote by $\mathcal{C}(\mathbb{Q}, \mathbb{P}) \subseteq \mathcal{P}(\Xi \times \Xi)$ the set of couplings of the distributions $\mathbb{Q} \in \mathcal{P}(\Xi)$ and $\mathbb{P} \in \mathcal{P}(\Xi)$.

Lemma 2. *Let $\mathbb{Q}, \mathbb{P}^1, \mathbb{P}^2 \in \mathcal{P}(\Xi)$ be probability distributions. If $\Pi^1 \in \mathcal{C}(\mathbb{Q}, \mathbb{P}^1)$ and $\Pi^2 \in \mathcal{C}(\mathbb{Q}, \mathbb{P}^2)$, then, $\lambda \cdot \Pi^1 + (1 - \lambda) \cdot \Pi^2 \in \mathcal{C}(\mathbb{Q}, \lambda \cdot \mathbb{P}^1 + (1 - \lambda) \cdot \mathbb{P}^2)$ for all $\lambda \in (0, 1)$.*

Proof. Let $\Pi = \lambda \cdot \Pi^1 + (1 - \lambda) \cdot \Pi^2$ and $\mathbb{P} = \lambda \cdot \mathbb{P}^1 + (1 - \lambda) \cdot \mathbb{P}^2$. To have $\Pi \in \mathcal{C}(\mathbb{Q}, \mathbb{P})$ we need $\Pi(d\xi, \Xi) = \mathbb{Q}(d\xi)$ and $\Pi(\Xi, d\xi') = \mathbb{P}(d\xi')$. To this end, observe that

$$\begin{aligned}
\Pi(d\xi, \Xi) &= \lambda \cdot \Pi^1(d\xi, \Xi) + (1 - \lambda) \cdot \Pi^2(d\xi, \Xi) \\
&= \lambda \cdot \mathbb{Q} + (1 - \lambda) \cdot \mathbb{Q} = \mathbb{Q}
\end{aligned}$$

where the second identity uses the fact that $\Pi^1 \in \mathcal{C}(\mathbb{Q}, \mathbb{P}^1)$. Similarly, we can show:

$$\begin{aligned}
\Pi(\Xi, d\xi') &= \lambda \cdot \Pi^1(\Xi, d\xi') + (1 - \lambda) \cdot \Pi^2(\Xi, d\xi') \\
&= \lambda \cdot \mathbb{P}^1 + (1 - \lambda) \cdot \mathbb{P}^2 = \mathbb{P},
\end{aligned}$$

which concludes the proof. $\square$

We further prove the following intermediary result.

Lemma 3. *Let $\mathbb{Q}, \mathbb{P}^1, \mathbb{P}^2 \in \mathcal{P}(\Xi)$ and $\mathbb{P} = \lambda \cdot \mathbb{P}^1 + (1 - \lambda) \cdot \mathbb{P}^2$ for some $\lambda \in (0, 1)$. We have:*

$$W(\mathbb{Q}, \mathbb{P}) \leq \lambda \cdot W(\mathbb{Q}, \mathbb{P}^1) + (1 - \lambda) \cdot W(\mathbb{Q}, \mathbb{P}^2).$$

Proof. The Wasserstein distance between $\mathbb{Q}, \mathbb{Q}' \in \mathcal{P}(\Xi)$ can be written as:

$$W(\mathbb{Q}, \mathbb{Q}') = \min_{\Pi \in \mathcal{C}(\mathbb{Q}, \mathbb{Q}')} \left\{ \int_{\Xi \times \Xi} d(\xi, \xi') \Pi(d\xi, d\xi') \right\},$$

and since d is a feature-label metric (cf. Definition 1) the minimum is well-defined (Villani, 2009, Theorem 4.1). We name the optimal solutions to the above problem the *optimal couplings*. Let Π^1 be an optimal coupling of $W(\mathbb{Q}, \mathbb{P}^1)$ and let Π^2 be an optimal coupling of $W(\mathbb{Q}, \mathbb{P}^2)$ and define $\Pi^c = \lambda \cdot \Pi^1 + (1 - \lambda) \cdot \Pi^2$. We have

$$\begin{aligned} W(\mathbb{Q}, \mathbb{P}) &= \min_{\Pi \in \mathcal{C}(\mathbb{Q}, \mathbb{P})} \left\{ \int_{\Xi \times \Xi} d(\xi, \xi') \Pi(d\xi, d\xi') \right\} \\ &\leq \int_{\Xi \times \Xi} d(\xi, \xi') \Pi^c(d\xi, d\xi') \\ &= \lambda \cdot \int_{\Xi \times \Xi} d(\xi, \xi') \Pi^1(d\xi, d\xi') + (1 - \lambda) \cdot \int_{\Xi \times \Xi} d(\xi, \xi') \Pi^2(d\xi, d\xi') \\ &= \lambda \cdot W(\mathbb{Q}, \mathbb{P}^1) + (1 - \lambda) \cdot W(\mathbb{Q}, \mathbb{P}^2), \end{aligned}$$

where the first identity uses the definition of the Wasserstein metric, the inequality is due to Lemma 2 as Π^c is a feasible coupling (not necessarily optimal), the equality that follows uses the definition of Π^c and the linearity of integrals, and the final identity uses the fact that Π^1 and Π^2 were constructed to be the optimal couplings. $\square$

We now prove the proposition (we refer to $\mathbb{Q}_{\text{mix}}$ in the statement of this lemma simply as $\mathbb{Q}$). To prove $\mathbb{Q} \in \mathfrak{B}_\varepsilon(\mathbb{P}_N) \cap \mathfrak{B}_{\hat{\varepsilon}}(\hat{\mathbb{P}}_{\hat{N}})$, it is sufficient to show that $W(\mathbb{P}_N, \mathbb{Q}) \leq \varepsilon$ and $W(\hat{\mathbb{P}}_{\hat{N}}, \mathbb{Q}) \leq \hat{\varepsilon}$ jointly hold. By using Lemma 3, we can derive the following inequalities:

$$\begin{aligned} W(\mathbb{P}_N, \mathbb{Q}) &\leq \lambda \cdot \underbrace{W(\mathbb{P}_N, \mathbb{P}_N)}_{=0} + (1 - \lambda) \cdot W(\mathbb{P}_N, \hat{\mathbb{P}}_{\hat{N}}) \\ W(\hat{\mathbb{P}}_{\hat{N}}, \mathbb{Q}) &\leq \lambda \cdot W(\mathbb{P}_N, \hat{\mathbb{P}}_{\hat{N}}) + (1 - \lambda) \cdot \underbrace{W(\hat{\mathbb{P}}_{\hat{N}}, \hat{\mathbb{P}}_{\hat{N}})}_{=0}. \end{aligned}$$

Therefore, sufficient conditions on $W(\mathbb{P}_N, \mathbb{Q}) \leq \varepsilon$ and $W(\hat{\mathbb{P}}_{\hat{N}}, \mathbb{Q}) \leq \hat{\varepsilon}$ would be:

$$\begin{cases} (1 - \lambda) \cdot W(\mathbb{P}_N, \hat{\mathbb{P}}_{\hat{N}}) \leq \varepsilon \\ \lambda \cdot W(\mathbb{P}_N, \hat{\mathbb{P}}_{\hat{N}}) \leq \hat{\varepsilon}. \end{cases}$$

Moreover, given that $\varepsilon + \hat{\varepsilon} \geq W(\mathbb{P}_N, \hat{\mathbb{P}}_{\hat{N}})$, the sufficient conditions further simplify to

$$\begin{cases} (1 - \lambda) \cdot \hat{\varepsilon} \leq \lambda \cdot \varepsilon \\ \lambda \cdot \varepsilon \leq (1 - \lambda) \cdot \hat{\varepsilon}. \end{cases} \iff \lambda \cdot \varepsilon = (1 - \lambda) \cdot \hat{\varepsilon},$$

which is implied when $\frac{\lambda}{1 - \lambda} = \frac{\hat{\varepsilon}}{\varepsilon}$, concluding the proof. $\square$

B.9 PROOF OF THEOREM 2

Since each result in the statement of this theorem is abridged, we will present these results sequentially as separate results. We review the existing literature to characterize $\mathfrak{B}_\varepsilon(\mathbb{P}_N)$, in a similar fashion with the results presented in (Selvi et al., 2022, Appendix A) for the logistic loss, by revising them to the adversarial loss whenever necessary. The N -fold product distribution of $\mathbb{P}^0$ from which the training set $\mathbb{P}_N$ is constructed is denoted below by $[\mathbb{P}^0]^N$.

Theorem 4. Assume there exist $a > 1$ and $A > 0$ such that $\mathbb{E}_{\mathbb{P}^0}[\exp(\|\xi\|^a)] \leq A$ for a norm $\|\cdot\|$ on $\mathbb{R}^n$. Then, there are constants $c_1, c_2 > 0$ that only depend on $\mathbb{P}^0$ through a , A , and n , such that $[\mathbb{P}^0]^N(\mathbb{P}^0 \in \mathfrak{B}_\varepsilon(\mathbb{P}_N)) \geq 1 - \eta$ holds for any confidence level $\eta \in (0, 1)$ as long as the Wasserstein ball radius satisfies the following optimal characterization

$$\varepsilon \geq \begin{cases} \left(\frac{\log(c_1/\eta)}{c_2 \cdot N} \right)^{1/\max\{n, 2\}} & \text{if } N \geq \frac{\log(c_1/\eta)}{c_2} \\ \left(\frac{\log(c_1/\eta)}{c_2 \cdot N} \right)^{1/a} & \text{otherwise.} \end{cases}$$

Proof. The statement follows from Theorem 18 of Kuhn et al. (2019). The presented decay rate $\mathcal{O}(N^{-1/n})$ of ε as N increases is optimal (Fournier and Guillin, 2015). $\square$

Now that we gave a confidence for the unknown radius ε satisfying $\mathbb{P}^0 \in \mathfrak{B}_\varepsilon(\mathbb{P}_N)$, we analyze the underlying optimization problems. Most of the theory is well-established for logistic loss function, and in the following we show that similar results follow for the adversarial loss function. For convenience, we state **DR-ARO** again by using the adversarial loss function as defined in Observation 1:

$$\begin{aligned} & \underset{\beta}{\text{minimize}} && \sup_{\mathbb{Q} \in \mathfrak{B}_\varepsilon(\mathbb{P}_N)} \mathbb{E}_{\mathbb{Q}}[\ell_{\beta}^{\alpha}(\mathbf{x}, y)] \\ & \text{subject to} && \beta \in \mathbb{R}^n. \end{aligned} \tag{DR-ARO}$$

Theorem 5. *If the assumptions of Theorem 4 are satisfied and ε is chosen as in the statement of Theorem 4, then*

$$[\mathbb{P}^0]^N \left(\mathbb{E}_{\mathbb{P}^0}[\ell_{\beta^*}^{\alpha}(\mathbf{x}, y)] \leq \sup_{\mathbb{Q} \in \mathfrak{B}_\varepsilon(\mathbb{P}_N)} \mathbb{E}_{\mathbb{Q}}[\ell_{\beta^*}^{\alpha}(\mathbf{x}, y)] \right) \geq 1 - \eta$$

holds for all $\eta \in (0, 1)$ and all optimizers β^ of **DR-ARO**.*

Proof. The statement follows from Theorem 19 of Kuhn et al. (2019) given that ℓ_{β}^{α} is a finite-valued continuous loss function. $\square$

Theorem 5 states that the optimal value of **DR-ARO** overestimates the true loss with arbitrarily high confidence $1 - \eta$. Despite the desired overestimation of the true loss, we show that **DR-ARO** is still asymptotically consistent if we restrict the set of admissible β to a bounded set¹.

Theorem 6. *If we restrict the hypotheses β to a bounded set $\mathcal{H} \subseteq \mathbb{R}^n$, and parameterize ε as ε_N to show its dependency to the sample size, then, under the assumptions of Theorem 4, we have*

$$\sup_{\mathbb{Q} \in \mathfrak{B}_{\varepsilon_N}(\mathbb{P}_N)} \mathbb{E}_{\mathbb{Q}}[\ell_{\beta^*}^{\alpha}(\mathbf{x}, y)] \xrightarrow{N \rightarrow \infty} \mathbb{E}_{\mathbb{P}^0}[\ell_{\beta^*}^{\alpha}(\mathbf{x}, y)] \quad \mathbb{P}^0\text{-almost surely,}$$

whenever ε_N is set as specified in Theorem 4 along with its finite-sample confidence η_N , and they satisfy $\sum_{N \in \mathbb{N}} \eta_N < \infty$ and $\lim_{N \rightarrow \infty} \varepsilon_N = 0$.

Proof. If we show that there exists $\xi^0 \in \Xi$ and $C > 0$ such that $\ell_{\beta}^{\alpha}(\mathbf{x}, y) \leq C(1 + d(\xi, \xi^0))$ holds for all $\beta \in \mathcal{H}$ and $\xi \in \Xi$ (that is, the adversarial loss satisfies a growth condition), the statement will follow immediately from Theorem 20 of (Kuhn et al., 2019).

To see that the growth condition is satisfied, we first substitute the definition of ℓ_{β}^{α} and d explicitly, and note that we would like to show there exists $\xi^0 \in \Xi$ and $C > 0$ such that

$$\log(1 + \exp(-y \cdot \beta^{\top} \mathbf{x} + \alpha \cdot \|\beta\|_{p^*})) \leq C(1 + \|\mathbf{x} - \mathbf{x}^0\|_q + \kappa \cdot \mathbb{1}[y \neq y^0])$$

holds for all $\beta \in \mathcal{H}$ and $\xi \in \Xi$. We take $\xi^0 = (\mathbf{0}, y^0)$ and show that the right-hand side of the inequality can be lower bounded as:

$$\begin{aligned} C(1 + \|\mathbf{x} - \mathbf{x}^0\|_q + \kappa \cdot \mathbb{1}[y \neq y^0]) &= C(1 + \|\mathbf{x}\|_q + \kappa \cdot \mathbb{1}[y \neq y^0]) \\ &\geq C(1 + \|\mathbf{x}\|_q). \end{aligned}$$

Moreover, the left-hand side of the inequality can be upper bounded for any $\beta \in \mathcal{H} \subseteq [-M, M]^n$ (for some $M > 0$) and

¹Note that, this is without loss of generality given that we can normalize the decision boundary of linear classifiers.

$\xi = (x, y) \in \Xi$ as:

$$\begin{aligned}
\log(1 + \exp(-y \cdot \beta^\top x + \alpha \cdot \|\beta\|_{p^*})) &\leq \log(1 + \exp(|\beta^\top x| + \alpha \cdot \|\beta\|_{p^*})) \\
&\leq \log(2 \cdot \exp(|\beta^\top x| + \alpha \cdot \|\beta\|_{p^*})) \\
&= \log(2) + |\beta^\top x| + \alpha \cdot \|\beta\|_{p^*} \\
&\leq \log(2) + \sup_{\beta \in [-M, M]^n} \{|\beta^\top x|\} + \alpha \cdot \sup_{\beta \in [-M, M]^n} \{\|\beta\|_{p^*}\} \\
&= \log(2) + M \cdot \|x\|_1 + M \cdot \alpha \\
&\leq \log(2) + M \cdot n^{(q-1)/q} \cdot \|x\|_1 + M \cdot \alpha
\end{aligned}$$

where the final inequality uses Hölder's inequality to bound the 1-norm with the q -norm. Thus, it suffices to show that we have

$$\log(2) + M \cdot n^{(q-1)/q} \cdot \|x\|_1 + M \cdot \alpha \leq C(1 + \|x\|_q) \quad \forall \xi \in \Xi,$$

which is satisfied for any $C \geq \max\{\log(2) + M \cdot \alpha, M \cdot n^{(q-1)/q}\}$. This completes the proof by showing the growth condition is satisfied. $\square$

So far, we reviewed tight characterizations for ε so that the ball $\mathfrak{B}_\varepsilon(\mathbb{P}_N)$ includes the true distribution $\mathbb{P}^0$ with arbitrarily high confidence, proved that the DRO problem **DR-ARO** overestimates the true loss, while converging to the true problem asymptotically as the confidence $1 - \eta$ increases and the radius ε decreases simultaneously. Finally, we discuss that for optimal solutions β^* to **DR-ARO**, there are worst case distributions $\mathbb{Q}^* \in \mathfrak{B}_\varepsilon(\mathbb{P}_N)$ of nature's problem that are supported on at most $N + 1$ outcomes.

Theorem 7. *If we restrict the hypotheses β to a bounded set $\mathcal{H} \subseteq \mathbb{R}^n$, then there are distributions $\mathbb{Q}^* \in \mathfrak{B}_\varepsilon(\mathbb{P}_N)$ that are supported on at most $N + 1$ outcomes and satisfy:*

$$\mathbb{E}_{\mathbb{Q}^*}[\ell_\beta^\alpha(x, y)] = \sup_{\mathbb{Q} \in \mathfrak{B}_\varepsilon(\mathbb{P}_N)} \mathbb{E}_{\mathbb{Q}}[\ell_\beta^\alpha(x, y)].$$

Proof. The proof follows from (Yue et al., 2022). $\square$

See the proof of Selvi et al. (2022, Theorem 8) and the discussion that follows for insights and further analysis on these results presented.

B.10 PROOF OF THEOREM 3

Firstly, since $\hat{\mathbb{P}}_{\hat{N}}$ is constructed from i.i.d. samples of $\hat{\mathbb{P}}$, we can overestimate the distance $\hat{\varepsilon}_1 = W(\hat{\mathbb{P}}_{\hat{N}}, \hat{\mathbb{P}})$ analogously by applying Theorem 4, *mutatis mutandis*. This leads us to the following result where the joint (independent) N -fold product distribution of $\mathbb{P}^0$ and the $\hat{N}$ -fold product distribution of $\hat{\mathbb{P}}$ is denoted below by $[\mathbb{P}^0 \times \hat{\mathbb{P}}]^{N \times \hat{N}}$.

Theorem 8. *Assume that there exist $a > 1$ and $A > 0$ such that $\mathbb{E}_{\mathbb{P}^0}[\exp(\|\xi\|^a)] \leq A$, and there exist $\hat{a} > 1$ and $\hat{A} > 0$ such that $\mathbb{E}_{\hat{\mathbb{P}}}[\exp(\|\xi\|^{\hat{a}})] \leq \hat{A}$ for a norm $\|\cdot\|$ on $\mathbb{R}^n$. Then, there are constants $c_1, c_2 > 0$ that only depends on $\mathbb{P}^0$ through a, A , and n , and constants $\hat{c}_1, \hat{c}_2 > 0$ that only depends on $\hat{\mathbb{P}}$ through $\hat{a}, \hat{A}$, and n such that $[\mathbb{P}^0 \times \hat{\mathbb{P}}]^{N \times \hat{N}}(\mathbb{P}^0 \in \mathfrak{B}_\varepsilon(\mathbb{P}_N) \cap \mathfrak{B}_{\hat{\varepsilon}}(\hat{\mathbb{P}}_{\hat{N}})) \geq 1 - \eta$ holds for any confidence level $\eta \in (0, 1)$ as long as the Wasserstein ball radii satisfy the following characterization*

$$\begin{aligned}
\varepsilon &\geq \begin{cases} \left(\frac{\log(c_1/\eta_1)}{c_2 \cdot N} \right)^{1/\max\{n, 2\}} & \text{if } N \geq \frac{\log(c_1/\eta_1)}{c_2} \\ \left(\frac{\log(c_1/\eta_1)}{c_2 \cdot N} \right)^{1/a} & \text{otherwise} \end{cases} \\
\hat{\varepsilon} &\geq W(\mathbb{P}^0, \hat{\mathbb{P}}) + \begin{cases} \left(\frac{\log(\hat{c}_1/\eta_2)}{\hat{c}_2 \cdot \hat{N}} \right)^{1/\max\{n, 2\}} & \text{if } \hat{N} \geq \frac{\log(\hat{c}_1/\eta_2)}{\hat{c}_2} \\ \left(\frac{\log(\hat{c}_1/\eta_2)}{\hat{c}_2 \cdot \hat{N}} \right)^{1/\hat{a}} & \text{otherwise} \end{cases}
\end{aligned}$$

for some $\eta_1, \eta_2 > 0$ satisfying $\eta_1 + \eta_2 = \eta$.

Proof. It immediately follows from Theorem 4 that $[\mathbb{P}^0]^N(\mathbb{P}^0 \in \mathfrak{B}_\varepsilon(\mathbb{P}_N)) \geq 1 - \eta_1$ holds. If we take $\hat{\varepsilon}_1 > 0$ as

$$\hat{\varepsilon}_1 \geq \begin{cases} \left(\frac{\log(\hat{c}_1/\eta_2)}{\hat{c}_2 \cdot \hat{N}} \right)^{1/\max\{n,2\}} & \text{if } \hat{N} \geq \frac{\log(\hat{c}_1/\eta_2)}{\hat{c}_2} \\ \left(\frac{\log(\hat{c}_1/\eta_2)}{\hat{c}_2 \cdot \hat{N}} \right)^{1/\hat{a}} & \text{otherwise,} \end{cases}$$

then, we similarly have $[\hat{\mathbb{P}}]^{\hat{N}}(\hat{\mathbb{P}} \in \mathfrak{B}_{\hat{\varepsilon}_1}(\hat{\mathbb{P}}_{\hat{N}})) \geq 1 - \eta_2$. Since the following implication follows from the triangle inequality:

$$\hat{\mathbb{P}} \in \mathfrak{B}_{\hat{\varepsilon}_1}(\hat{\mathbb{P}}_{\hat{N}}) \implies \mathbb{P}^0 \in \mathfrak{B}_{\hat{\varepsilon}_1 + \mathbf{W}(\mathbb{P}^0, \hat{\mathbb{P}})}(\hat{\mathbb{P}}_{\hat{N}}),$$

we have that $[\hat{\mathbb{P}}]^{\hat{N}}(\mathbb{P}^0 \in \mathfrak{B}_\varepsilon(\hat{\mathbb{P}}_{\hat{N}})) \geq 1 - \eta_2$. These results, along with the facts that $\hat{\mathbb{P}}_{\hat{N}}$ and $\mathbb{P}_N$ are independently sampled from their true distributions, imply:

$$\begin{aligned} & [\mathbb{P}^0 \times \hat{\mathbb{P}}]^{N \times \hat{N}}(\mathbb{P}^0 \notin \mathfrak{B}_\varepsilon(\mathbb{P}_N) \vee \mathbb{P}^0 \notin \mathfrak{B}_{\hat{\varepsilon}}(\hat{\mathbb{P}}_{\hat{N}})) \\ & \leq [\mathbb{P}^0 \times \hat{\mathbb{P}}]^{N \times \hat{N}}(\mathbb{P}^0 \notin \mathfrak{B}_\varepsilon(\mathbb{P}_N)) + [\mathbb{P}^0 \times \hat{\mathbb{P}}]^{N \times \hat{N}}(\mathbb{P}^0 \notin \mathfrak{B}_{\hat{\varepsilon}}(\hat{\mathbb{P}}_{\hat{N}})) \\ & = [\mathbb{P}^0]^N(\mathbb{P}^0 \notin \mathfrak{B}_\varepsilon(\mathbb{P}_N)) + [\hat{\mathbb{P}}]^{\hat{N}}(\mathbb{P}^0 \notin \mathfrak{B}_{\hat{\varepsilon}}(\hat{\mathbb{P}}_{\hat{N}})) < \eta_1 + \eta_2 \end{aligned}$$

implying the desired result $[\mathbb{P}^0 \times \hat{\mathbb{P}}]^{N \times \hat{N}}(\mathbb{P}^0 \in \mathfrak{B}_\varepsilon(\mathbb{P}_N) \cap \mathfrak{B}_{\hat{\varepsilon}}(\hat{\mathbb{P}}_{\hat{N}})) \geq 1 - \eta$. $\square$

The second statement immediately follows under the assumptions of Theorem 8: **Inter-ARO** overestimates the true loss analogously as Theorem 5 with an identical proof.

C EXPONENTIAL CONIC REFORMULATION OF **DR-ARO**

For any $i \in [N]$, the constraints of **DR-ARO** are

$$\begin{cases} \log(1 + \exp(-y^i \cdot \boldsymbol{\beta}^\top \mathbf{x}^i + \alpha \cdot \|\boldsymbol{\beta}\|_{p^*})) \leq s_i \\ \log(1 + \exp(y^i \cdot \boldsymbol{\beta}^\top \mathbf{x}^i + \alpha \cdot \|\boldsymbol{\beta}\|_{p^*})) - \lambda \cdot \kappa \leq s_i, \end{cases}$$

which, by using an auxiliary variable u , can be written as

$$\begin{cases} \log(1 + \exp(-y^i \cdot \boldsymbol{\beta}^\top \mathbf{x}^i + u)) \leq s_i \\ \log(1 + \exp(y^i \cdot \boldsymbol{\beta}^\top \mathbf{x}^i + u)) - \lambda \cdot \kappa \leq s_i \\ \alpha \cdot \|\boldsymbol{\beta}\|_{p^*} \leq u. \end{cases}$$

Following the conic modeling guidelines of [MOSEK ApS \(2023\)](#), for new variables $v_i^+, w_i^+ \in \mathbb{R}$, the first constraint can be written as

$$\{v_i^+ + w_i^+ \leq 1, (v_i^+, 1, [-u + y^i \cdot \boldsymbol{\beta}^\top \mathbf{x}^i] - s_i) \in \mathcal{K}_{\text{exp}}, (w_i^+, 1, -s_i) \in \mathcal{K}_{\text{exp}},$$

by using the definition of the exponential cone $\mathcal{K}_{\text{exp}}$. Similarly, for new variables $v_i^-, w_i^- \in \mathbb{R}$, the second constraint can be written as

$$\{v_i^- + w_i^- \leq 1, (v_i^-, 1, [-u - y^i \cdot \boldsymbol{\beta}^\top \mathbf{x}^i] - s_i - \lambda \cdot \kappa) \in \mathcal{K}_{\text{exp}}, (w_i^-, 1, -s_i - \lambda \cdot \kappa) \in \mathcal{K}_{\text{exp}}.$$

Applying this for all $i \in [N]$ concludes that the following is the conic formulation of **DR-ARO**:

$$\begin{aligned}
& \underset{\beta, \lambda, \mathbf{s}, u}{\text{minimize}} && \lambda \cdot \varepsilon + \frac{1}{N} \sum_{i \in [N]} s_i \\
& \underset{\mathbf{v}^+, \mathbf{w}^+, \mathbf{v}^-, \mathbf{w}^-}{\text{subject to}} && v_i^+ + w_i^+ \leq 1 && \forall i \in [N] \\
& && (v_i^+, 1, [-u + y^i \cdot \beta^\top \mathbf{x}^i] - s_i) \in \mathcal{K}_{\text{exp}}, (w_i^+, 1, -s_i) \in \mathcal{K}_{\text{exp}} && \forall i \in [N] \\
& && v_i^- + w_i^- \leq 1 && \forall i \in [N] \\
& && (v_i^-, 1, [-u - y^i \cdot \beta^\top \mathbf{x}^i] - s_i - \lambda \cdot \kappa) \in \mathcal{K}_{\text{exp}}, (w_i^-, 1, -s_i - \lambda \cdot \kappa) \in \mathcal{K}_{\text{exp}} && \forall i \in [N] \\
& && \alpha \cdot \|\beta\|_{p^*} \leq u \\
& && \|\beta\|_{q^*} \leq \lambda \\
& && \beta \in \mathbb{R}^n, \lambda \geq 0, \mathbf{s} \in \mathbb{R}^N, u \in \mathbb{R}, \mathbf{v}^+, \mathbf{w}^+, \mathbf{v}^-, \mathbf{w}^- \in \mathbb{R}^N.
\end{aligned}$$

D FURTHER DETAILS ON NUMERICAL EXPERIMENTS

D.1 UCI EXPERIMENTS

Preprocessing UCI datasets. We experiment on 10 UCI datasets (Kelly et al., 2023) (cf. Table 3). We use Python 3 for preprocessing these datasets. Classification problems with more than two classes are converted to binary classification problems (most frequent class/others). For all datasets, numerical features are standardized, the ordinal categorical features are left as they are, and the nominal categorical features are processed via one-hot encoding. As mentioned in the main paper, we obtain auxiliary (synthetic) datasets via SDV, which is also implemented in Python 3.

Table 3: Size of the UCI datasets.

DataSet	N	$\hat{N}$	N_{te}	n
absent	111	333	296	74
annealing	134	404	360	41
audiology	33	102	91	102
breast-cancer	102	307	274	90
contraceptive	220	663	590	23
dermatology	53	161	144	99
ecoli	50	151	135	9
spambase	690	2,070	1,841	58
spect	24	72	64	23
prim-tumor	50	153	136	32

Detailed misclassification results on the UCI datasets. Table 4 contains detailed results on the out-of-sample error rates of each method on 10 UCI datasets for classification. All parameters are 5-fold cross-validated: Wasserstein radii from the grid $\{10^{-6}, 10^{-5}, 10^{-4}, 10^{-3}, 10^{-2}, 10^{-1}, 0, 1, 2, 5, 10\}$, κ from the grid $\{1, \sqrt{n}, n\}$ the weight parameter of ARO+Aux from grid $\{10^{-6}, 10^{-5}, 10^{-4}, 10^{-3}, 10^{-2}, 10^{-1}, 0, 1\}$. We fix the norm defining the feature-label metric to the ℓ_1 -norm, and test ℓ_2 -attacks, but other choices with analogous results are also implemented.

Finally, we demonstrate that our theory, especially DRO+ARO+Aux, contributes to the DRO literature even without adversarial attacks. In this case of $\alpha = 0$, ERM and ARO would be equivalent, and DRO+ARO would reduce to the traditional DR LR model (Shafieezadeh-Abadeh et al., 2015). ARO+Aux would be interpreted as revising the empirical distribution of ERM to a mixture (mixture weight cross-validated) of the empirical and auxiliary distributions. DRO+ARO+Aux, on the other hand, can be interpreted as DRO over a carefully reduced ambiguity set (intersection of the empirical and auxiliary Wasserstein balls). The results are in Table 5. Analogous results follow as before (that is, DRO+ARO+Aux is the ‘winning’ approach, DRO+ARO and ARO+Aux alternate for the ‘second’ approach), with the exception of the dataset *contraceptive*, where ARO+Aux outperforms others.

Table 4: Mean ($\pm$ std) out-of-sample errors of UCI datasets, each with 10 simulations. Results for adversarial (ℓ_2 -)attack strengths $\alpha = 0.05$ and $\alpha = 0.2$ are shared.

Data	α	ERM	ARO	ARO+Aux	DRO+ARO	DRO+ARO+Aux
absent	0.05	44.02% (± 2.89)	38.82% (± 2.86)	35.95% (± 3.78)	34.22% (± 2.70)	32.64% (± 2.54)
	0.20	73.65% (± 4.14)	51.49% (± 3.39)	49.56% (± 3.80)	45.61% (± 2.32)	44.90% (± 2.30)
annealing	0.05	18.08% (± 1.89)	16.61% (± 2.16)	14.97% (± 1.39)	13.50% (± 2.98)	12.78% (± 2.78)
	0.20	37.31% (± 3.92)	23.08% (± 2.82)	21.30% (± 1.93)	20.70% (± 1.32)	19.53% (± 1.42)
audiology	0.05	21.43% (± 3.64)	21.54% (± 3.92)	17.03% (± 2.90)	11.76% (± 3.28)	9.01% (± 3.54)
	0.20	37.91% (± 6.78)	29.34% (± 5.89)	20.44% (± 2.75)	20.00% (± 3.01)	17.91% (± 3.28)
breast-cancer	0.05	4.74% (± 1.26)	4.93% (± 1.75)	3.87% (± 1.17)	3.06% (± 0.79)	2.52% (± 0.50)
	0.20	9.93% (± 1.73)	8.14% (± 2.01)	6.09% (± 1.79)	5.04% (± 1.11)	4.67% (± 0.99)
contraceptive	0.05	44.14% (± 2.80)	42.86% (± 2.59)	40.98% (± 0.95)	40.00% (± 1.33)	39.65% (± 1.15)
	0.20	66.19% (± 5.97)	43.49% (± 2.24)	42.71% (± 1.47)	42.71% (± 1.47)	42.71% (± 1.47)
dermatology	0.05	15.97% (± 2.64)	16.46% (± 1.67)	13.47% (± 1.97)	12.78% (± 1.61)	10.84% (± 1.24)
	0.20	30.07% (± 4.24)	28.54% (± 3.25)	21.53% (± 2.17)	22.64% (± 2.15)	20.21% (± 1.58)
ecoli	0.05	16.30% (± 4.42)	14.67% (± 5.13)	13.26% (± 3.07)	11.11% (± 5.52)	9.78% (± 2.61)
	0.20	51.41% (± 3.37)	42.67% (± 2.91)	41.85% (± 2.95)	39.70% (± 2.68)	38.89% (± 2.57)
spambase	0.05	11.35% (± 0.77)	10.23% (± 0.54)	10.16% (± 0.56)	9.83% (± 0.37)	9.81% (± 0.38)
	0.20	27.32% (± 2.11)	15.83% (± 0.77)	15.70% (± 0.76)	15.67% (± 0.72)	15.50% (± 0.68)
spect	0.05	33.75% (± 5.17)	29.69% (± 5.46)	25.78% (± 3.06)	25.47% (± 3.38)	21.56% (± 2.74)
	0.20	54.22% (± 9.88)	37.5% (± 3.53)	35.16% (± 2.47)	33.75% (± 2.68)	30.16% (± 3.61)
prim-tumor	0.05	21.84% (± 4.55)	20.81% (± 3.97)	17.35% (± 3.59)	16.18% (± 3.83)	14.78% (± 2.89)
	0.20	34.19% (± 6.17)	25.37% (± 4.58)	21.62% (± 3.45)	21.84% (± 3.34)	19.63% (± 2.71)

Table 5: Mean out-of-sample errors of UCI experiments without adversarial attacks.

Data	ERM	ARO	ARO+Aux	DRO+ARO	DRO+ARO+Aux
absent	36.28%	36.28%	31.86%	28.31%	27.74%
annealing	10.61%	10.61%	7.64%	7.14%	7.14%
audiology	14.94%	14.94%	12.97%	10.11%	7.69%
breast-cancer	6.64%	6.64%	5.22%	2.55%	2.15%
contraceptive	35.00%	35.00%	33.75%	34.56%	33.85%
dermatology	16.04%	16.04%	11.60%	9.93%	8.06%
ecoli	6.74%	6.74%	4.96%	5.19%	4.37%
spambase	8.95%	8.95%	8.52%	8.34%	8.16%
spect	30.74%	30.74%	24.69%	22.35%	18.75%
prim-tumor	22.79%	22.79%	17.28%	15.07%	13.97%

Cross-validated Wasserstein radii. In Corollary 2, we discussed the feasibility of ignoring the auxiliary data when its data-generating distribution is distant from the true data-generating distribution. To examine whether large values of $\hat{\epsilon}$ are selected via cross-validation in our experiments, we investigated their histograms and observed that this indeed occurs frequently. For example, as shown in Table 5, when there are no adversarial attacks, DRO+ARO and DRO+ARO+Aux yield identical errors on the ‘annealing’ dataset. This is because the two models become equivalent: DRO+ARO+Aux selects a large $\hat{\epsilon}$ (either 5 or 10), which ensures that the intersection of the Wasserstein balls reduces to the empirical distribution. With $\alpha = 0.05$, DRO+ARO+Aux selects a large $\hat{\epsilon}$ in 7 out of 10 simulations, while in the remaining 3 it selects a smaller value, resulting in a nontrivial intersection and, in turn, improved performance over DRO+ARO. We conclude that our method selects a nontrivial $\hat{\epsilon}$ only when there is evidence that the auxiliary data is *useful*, that is, when its data-generating distribution is sufficiently close to the true one. For instance, on the ‘absent’ dataset, we always have $\hat{\epsilon} \in \{10^{-2}, 10^{-1}\}$. We also revised our numerical experiments by modifying the Gaussian copula synthesizer to enforce uniform marginals (which results in significant information loss), and in this setting, our models consistently selected large $\hat{\epsilon}$ values.

Different training/auxiliary dataset ratio. Recall that in the UCI experiments, we sampled 15% of the original dataset as the training set, and used 45% to generate a synthetic auxiliary dataset. This setup was chosen to simulate scenarios

Table 6: Additional MNIST/EMNIST Benchmark.

Attack	ERM	ARO	ARO+Aux	DRO++	DRO+ARO	DRO+ARO+Aux
No attack ($\alpha = 0$)	1.55%	1.55%	1.19%	0.72%	0.64%	0.53%
ℓ_1 ($\alpha = 68/255$)	2.17%	1.84%	1.33%	1.40%	0.66%	0.57%
ℓ_2 ($\alpha = 128/255$)	99.93%	3.36%	2.54%	2.72%	2.40%	2.12%
ℓ_∞ ($\alpha = 8/255$)	100.00%	2.60%	2.38%	2.31%	2.20%	1.95%

with limited training data, where overfitting is a particular concern. If we reverse these proportions and use 45% of the original dataset for training and 15% for generating auxiliary data instead, we find that the best-performing method remains DRO+ARO+Aux, and the relative ranking among the models remains qualitatively similar. One notable difference is that the improvement of the auxiliary-data-oblivious DRO+ARO method over the non-DR ARO+Aux model becomes less significant, and is even reversed in the case of $\alpha = 0.05$ on the ‘ecoli’ dataset: ERM (14.59%), ARO (11.41%), ARO+Aux (9.04%), DRO+ARO (9.41%), DRO+ARO+Aux (8.89%). As expected, since more data is drawn from the true data-generating distribution, all methods exhibit improved out-of-sample performance compared to the original setup. However, in this case, DRO+ARO no longer outperforms ARO+Aux.

D.2 MNIST/EMNIST EXPERIMENTS

The setting in the MNIST/EMNIST experiments is similar to that in the UCI experiments. However, for auxiliary data, we use the EMNIST dataset which we accessed via the MLDatasets package of Julia.

Moreover, in §2 we reviewed the literature showing that when statistical error is not a concern, that is, when optimizing over the empirical distribution cannot cause overfitting (*e.g.*, in high-data regimes), then adversarial training is equivalent to a type- ∞ Wasserstein DRO problem with radius $\varepsilon = \alpha$. Hence, a natural question is whether increasing the value of ε further also provides distributional robustness. To this end, in Table 6, we revise Table 2 and add an additional benchmark DRO++. Here, we take $\varepsilon = \alpha + \varepsilon'$, and cross-validate ε' from the same grid that we cross-validated ε for methods DRO+ARO and DRO+ARO+Aux. We observe that DRO++ does not improve over DRO+ARO, which is expected given that type- ∞ Wasserstein DRO does not provide better generalizations, unlike type-1 Wasserstein DRO (*cf.* the discussion at the end of §3). Yet, this method improves over ARO in all cases, and even over ARO+Aux in the no-attack or ℓ_∞ -attack settings.

D.3 ARTIFICIAL EXPERIMENTS

Data generation. We sample a ‘true’ β from a unit ℓ_2 -ball, and generate data as summarized in Algorithm 1. Such a dataset generation gives N instances from the same true data-generating distribution. In order to obtain N auxiliary dataset instances, we perturb the probabilities p^i with standard random normal noise which is equivalent to sampling i.i.d. from a *perturbed* distribution. Testing is always done on true data, that is, the test set is sampled according to Algorithm 1.

Algorithm 1 Data from a ground truth logistic classifier

Input: set of feature vectors $\mathbf{x}^i$, $i \in [N]$; vector β

for $i \in \{1, \dots, N\}$ **do**

Find the probability $p^i = [1 + \exp(-\beta^\top \mathbf{x}^i)]^{-1}$.

Sample $u = \mathcal{U}(0, 1)$

if $p^i \geq u$ **then**

$y^i = +1$

else

$y^i = -1$

end if

end for

Output: $(\mathbf{x}^i, y^i)$, $i \in [N]$.

Table 7: Mean w in problem (2) and $\varepsilon/\hat{\varepsilon}$ in problem **Inter-ARO** across 25 simulations of cross-validating ω , ε , and $\hat{\varepsilon}$.

Attack	ARO+Aux (cross-validated w)	DRO+ARO+Aux (cross-validated $\varepsilon/\hat{\varepsilon}$)
$\alpha = 0$	0.002	0.0120
$\alpha = 0.1$	0.046	0.172
$\alpha = 0.25$	0.086	0.232
$\alpha = 0.5$	0.290	0.241

Strength of the attack and importance of auxiliary data. In the main paper we discussed how the strength of an attack determines whether using auxiliary data in ARO (ARO+Aux) or considering distributional ambiguity (DRO+ARO) is more effective, and observed that unifying them to obtain DRO+ARO+Aux yields the best results in all attack regimes. Now we focus on the methods that rely on auxiliary data, namely ARO+Aux and DRO+ARO+Aux and explore the importance of auxiliary data $\hat{\mathbb{P}}_{\hat{N}}$ in comparison to its empirical counterpart $\mathbb{P}_N$. Table 7 shows the average values of w for problem (2) obtained via cross-validation. We see that the greater the attack strength is the more we should use the auxiliary data in ARO+Aux. The same relationship holds for the average of $\varepsilon/\hat{\varepsilon}$ obtained via cross-validation in **Inter-ARO**, which means that the relative size of the Wasserstein ball built around the empirical distribution gets larger compared to the same ball around the auxiliary data, that is, ambiguity around the auxiliary data is smaller than the ambiguity around the empirical data. We highlight as a possible future research direction exploring when a larger attack *per se* implies the intersection will move towards the auxiliary data distribution.

More results on scalability. We further simulate 25 cases with an ℓ_2 -attack strength of $\alpha = 0.2$, $N = 200$ instances in the training dataset, $\hat{N} = 200$ instances in the auxiliary dataset, and we vary the number of features n . We report the median ($50\% \pm 15\%$ quantiles shaded) runtimes of each method in Figure 3. The fastest methods are ERM and ARO among which the faster one depends on n (as the adversarial loss includes a regularizer of β), followed by ARO+Aux, DRO+ARO, and DRO+ARO+Aux, respectively. DRO+ARO+Aux is the slowest, which is expected given that DRO+ARO is its special for large $\hat{\varepsilon}$. The runtime however scales gracefully.

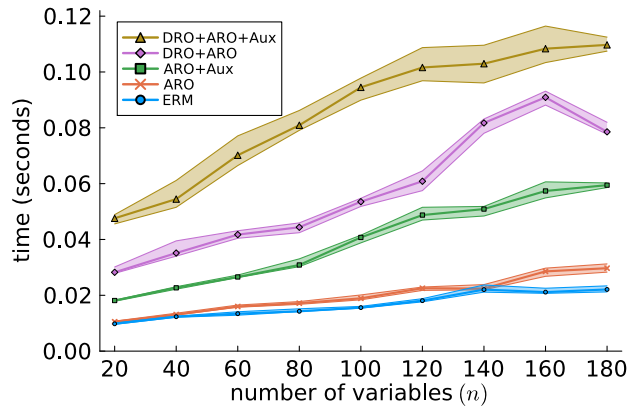


Figure 3: Runtimes under a varying number of features in the artificially generated empirical and auxiliary datasets.

Finally, we focus further on DRO+ARO+Aux which solves problem **Inter-ARO** with $\mathcal{O}(n \cdot N \cdot \hat{N})$ variables and exponential cone constraints. For $n = 1,000$ and $N = \hat{N} = 10,000$, we observe that the runtimes vary between 134 to 232 seconds across 25 simulations.