

MITIGATING SPURIOUS CORRELATION VIA DISTRIBUTIONALLY ROBUST LEARNING WITH HIERARCHICAL AMBIGUITY SETS

Sung Ho Jo

Pohang University of Science and Technology
tjdgh1813@postech.ac.kr

Seonghwi Kim

Pohang University of Science and Technology
kshwi@postech.ac.kr

Minwoo Chae*

Pohang University of Science and Technology
mchae@postech.ac.kr

ABSTRACT

Conventional supervised learning methods are often vulnerable to spurious correlations, particularly under distribution shifts in test data. To address this issue, several approaches, most notably Group DRO, have been developed. While these methods are highly robust to subpopulation or group shifts, they remain vulnerable to intra-group distributional shifts, which frequently occur in minority groups with limited samples. We propose a hierarchical extension of Group DRO that addresses both inter-group and intra-group uncertainties, providing robustness to distribution shifts at multiple levels. We also introduce new benchmark settings that simulate realistic minority group distribution shifts—an important yet previously underexplored challenge in spurious correlation research. Our method demonstrates strong robustness under these conditions—where existing robust learning methods consistently fail—while also achieving superior performance on standard benchmarks. These results highlight the importance of broadening the ambiguity set to better capture both inter-group and intra-group distributional uncertainties.

1 INTRODUCTION

In recent years, machine learning methods have achieved remarkable success across a wide range of applications. An important objective of many machine learning methods is to learn model parameters that minimize the population risk, which is the population expectation of the loss function. Given training data and model parameters, the population risk can be approximated by the empirical risk, defined as the sample-averaged loss. Therefore, model parameters can be learned by minimizing the empirical risk, which is known as the empirical risk minimization (ERM) principle.

The underlying assumption of ERM-based methods is that the unseen future data, often referred to as test data, share the same distribution as the training data. However, in many real-world problems, the test data may follow a different distribution from the training data for various reasons. A notable example is subpopulation shift, where the training population consists of several groups (subpopulations), and the proportion of each group in the test data differs from that in the training data (Sagawa et al., 2020; Cai et al., 2021; Yang et al., 2023).

In many instances of subpopulation shifts, the group indicator is spuriously correlated with the target label or response variable. For example, in the widely studied Waterbirds dataset (Sagawa et al., 2020), the target label (e.g., “Waterbird” or “Landbird”) is spuriously correlated with the background environment (e.g., water or land). As a result, ERM-based models tend to associate “Waterbird” primarily with water backgrounds, leading to significant performance degradation on minority groups, such as waterbirds appearing against land backgrounds. These vulnerabilities extend beyond controlled benchmarks, posing substantial risks in real-world, high-stakes domain such as healthcare

*Corresponding author.

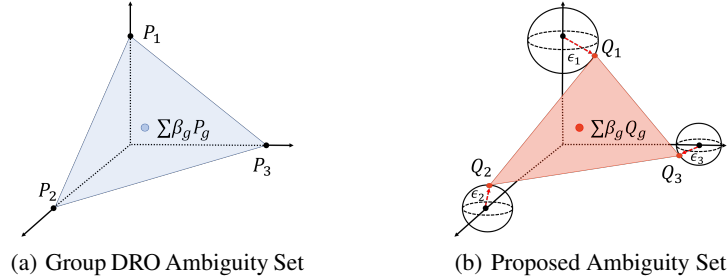


Figure 1: Comparison of the Group DRO ambiguity set (a) and our hierarchical extension (b). While Group DRO restricts uncertainty to mixtures of group distributions, our approach introduces additional within-group uncertainty (indicated by red dashed arrows), offering robustness to both inter-group and intra-group distributional shifts. (For visualization, we assume the 3-dimensional space in the figure represents a probability space, where each point corresponds to a probability distribution.)

(Zech et al., 2018; Badgeley et al., 2019), fairness (Hashimoto et al., 2018; Buolamwini & Gebru, 2018; Obermeyer et al., 2019), and autonomous driving (Zhang et al., 2017).

Over the past few years, a growing body of research has focused on mitigating these spurious correlations to ensure more reliable and stable model performance. Among the proposed methodologies, group distributionally robust optimization (Group DRO) (Sagawa et al., 2020) has emerged as a foundational approach. By partitioning the data into predefined groups and optimizing for the worst group loss, Group DRO effectively minimizes the model’s reliance on spurious correlations tied to specific subsets of data. This framework has inspired a wide range of subsequent methods, such as JTT (Liu et al., 2021), SSA (Nam et al., 2022), PG-DRO (Ghosal & Li, 2023), DISC (Wu et al., 2023) and GIC (Han & Zou, 2024), which commonly employ a two-step strategy: first identifying latent groups and then utilizing established robust training approaches—frequently Group DRO itself—to enhance model robustness.

While these methods have significantly advanced the field, they primarily focus on minimizing the worst-group loss under the assumption that each group’s training distribution reliably represents its true underlying distribution. However, this assumption often fails in practice—especially for minority groups with limited samples—where within-group distributional shifts naturally arise as a consequence of underrepresentation in the training data (Ben-Tal & Nemirovski, 2002; Devroye et al., 2013; Duchi & Namkoong, 2021). This limitation highlights the need for a more flexible and robust approach that accounts for uncertainty not only across groups, but also within them.

In this work, we address these limitations by introducing a hierarchical ambiguity set within the Group DRO framework, capturing both inter-group and intra-group uncertainties. As illustrated in Figure 1, while conventional Group DRO focuses on robustness only to shifts in group proportions by minimizing the worst-group risk, our approach extends this perspective by additionally modeling uncertainty within individual groups.

Technically, we employ a Wasserstein-distance-based formulation, which has recently garnered significant theoretical and empirical support for its efficacy in designing distributionally robust learning methods (Volpi et al., 2018; Kuhn et al., 2019; Blanchet et al., 2022; Bai et al., 2024). By defining a semantically meaningful cost function in a latent space, this formulation flexibly accommodates variations in the underlying data-generating mechanisms within each group. Consequently, our hierarchical ambiguity set enables the model to maintain robustness across a broader spectrum of distributional deviations, particularly for minority groups that are underrepresented in the training data.

Our main contributions are threefold:

- We propose a novel hierarchical ambiguity set that simultaneously captures inter-group and intra-group uncertainties, constituting a fundamental advancement over previous methods built on Group DRO.

- We develop a tractable minimax optimization algorithm that is computationally efficient, enabling the practical realization of our DRO framework with a hierarchical ambiguity set.
- We introduce new evaluation settings on standard benchmarks that reveal a previously overlooked yet important failure mode: minority-group distribution shifts—even when simply altering the train–test split without synthetic noise or external shifts—where state-of-the-art methods collapse but our approach achieves consistently strong performance.

2 RELATED WORK

Using explicit group labels is a well-established approach to achieving robust performance on underrepresented subpopulations. Sagawa et al. (2020) pioneered partitioning data by known group annotations and minimizing worst-group loss. Subsequent works extend this paradigm in multiple directions: Along similar lines, Krueger et al. (2021) expands Group DRO’s ambiguity set from convex to affine combinations of group distributions, although it still assumes fixed conditional distributions within each group. Meanwhile, Kirichenko et al. (2023) rebalances distributions via subsampling, Deng et al. (2024) employs a staged expansion of a group-balanced set, and Yao et al. (2022) uses Mixup-based augmentations (e.g., CutMix, Manifold Mix) to learn more invariant features. Izmailov et al. (2022) shows that even standard ERM can yield robust feature representations when combined with selective reweighting, and Piratla et al. (2022) further leverages inter-group interactions to identify shared features that enhance distributional robustness.

In parallel, a growing body of work addresses scenarios where group annotations are unavailable or prohibitively expensive, motivating the need to infer or approximate group structure. Building on Kirichenko et al. (2023), LaBonte et al. (2024) extends last-layer retraining to settings with minimal or no group annotations. Some methods employ limited validation sets to discover latent groups and then apply Group DRO (Sagawa et al., 2020) for robust optimization (Nam et al., 2022; Ghosal & Li, 2023). Indeed, the model’s loss (or its alternatives) often helps identify underrepresented subpopulations; for example, Nam et al. (2020); Liu et al. (2021); Qiu et al. (2023) use high-loss samples to recognize minority groups. Other approaches draw on diverse cues: Ahmed et al. (2021); Crea-ger et al. (2021) infer groups by maximizing violations of the invariant risk minimization principle (Arjovsky et al., 2019), while Asgari et al. (2022) employs masking to reduce reliance on spurious features. Sohoni et al. (2020); Seo et al. (2022) instead cluster feature embeddings to discover latent groups and identify pseudo-attributes for debiasing. Likewise, Zhang et al. (2022) proposes a two-stage contrastive learning framework by aligning samples with the same class but different spurious attributes, and Wu et al. (2023) constructs a “concept bank” of candidate spurious attributes for robust partitioning. Another line of work identifies spurious features by comparing the training set with a carefully selected reference dataset (Han & Zou, 2024) or removes examples that disproportionately degrade worst-group accuracy (Jain et al., 2024).

Beyond these established techniques, emerging efforts explore new failure modes within the spurious correlation framework. For instance, recent works have examined imperfect group partitions (Zhou et al., 2021), multiple spurious features (Paranjape et al., 2023; Kim et al., 2024), low spurious signal strength (Mulchandani & Kim, 2025), and spurious-aware out-of-distribution detection (Zohrabi et al., 2025). Relatedly, recent work has shown that the presence of spurious correlations can itself constitute a failure mode for domain adaptation methods (Kim et al., 2026). In this context, our work addresses a critical yet underexplored issue: minority-group data, by nature extremely limited in size, may not faithfully reflect their underlying distributions—a natural but overlooked phenomenon that must be explicitly accounted for.

3 PRELIMINARIES

Problem Setup. We consider a supervised learning problem where each observation consists of input features $X \in \mathcal{X}$ and a label $Y \in \mathcal{Y}$. Let $\{(x_i, y_i)\}_{i=1}^n$ be the training data and Θ be the parameter space. Assuming that the test data follows the same distribution as the training data, an important goal of various machine learning methods is to minimize the risk (also referred to as the test or generalization error)

$$\mathbb{E}_P[\mathcal{L}(f^\theta(X), Y)] \quad (1)$$

over Θ , where f^θ is a function parametrized by θ , $\mathcal{L}$ is a standard loss function, such as the cross-entropy, and $\mathbb{E}_P$ denotes the expectation under the population distribution P . To achieve this, one may solve the following optimization problem:

$$\inf_{\theta \in \Theta} \left\{ \mathbb{E}_{\hat{P}}[\ell(\theta; (X, Y))] = \frac{1}{n} \sum_{i=1}^n \ell(\theta; (x_i, y_i)) \right\},$$

often referred to as the ERM problem, where $\hat{P}$ is the empirical measure and $\ell(\theta; (X, Y))$ is shorthand for $\mathcal{L}(f^\theta(X), Y)$.

In addition to the above basic setting, we assume that the data are partitioned into multiple groups, with an indicator variable $G \in \mathcal{G}$; thus, the training data can be expressed as $\{(x_i, y_i, g_i)\}_{i=1}^n$. In particular, we assume that this indicator variable is spuriously correlated with the label Y . More specifically, in all our examples, we assume the existence of a spurious attribute $A \in \mathcal{A}$, and that the group variable $G = (Y, A)$ is a pair involving the response variable Y . Hence, $\mathcal{G} = \mathcal{Y} \times \mathcal{A}$. With adjusted notation, we write $\mathcal{G} = \{1, \dots, m\}$.

In the Waterbirds dataset, for example, the task is to classify birds as Waterbird or Landbird, with a spurious attribute being the background type (Water Background or Land Background), which results in four distinct groups. This dataset provides a classic example of spurious correlation: most waterbirds ($Y = \text{Waterbird}$) are found against water backgrounds ($A = \text{Water Background}$), leading models to rely excessively on background features. This overreliance substantially diminishes the performance of ERM on underrepresented groups, such as waterbirds on land backgrounds.

Group DRO. Group DRO (Sagawa et al., 2020) was devised to address the aforementioned issue caused by the spurious attribute. In the Group DRO, the training data are modeled as instances from a mixture distribution $P = \sum_{g=1}^m \alpha_g P_g$, where P_g denotes the conditional distribution of (X, Y) given $G = g$, and $\alpha = (\alpha_1, \dots, \alpha_m)$ represents the mixing proportion. Thus, each group forms a subpopulation within the training data. Instead of minimizing the population risk (1), Group DRO aims to minimize

$$\inf_{\theta \in \Theta} \max_{g \in \mathcal{G}} \mathbb{E}_{P_g}[\ell(\theta; (X, Y))], \quad (2)$$

which corresponds to the risk of the worst-performing group. Consequently, this procedure is highly robust to the subpopulation shift described in the introduction.

The Group DRO formulation (2) involves optimization over the discrete group variable g , posing a computational challenge for practical use. Sagawa et al. (2020) showed that the problem can be reformulated into an equivalent form

$$\inf_{\theta \in \Theta} \sup_{\beta \in \Delta_{m-1}} \sum_{g=1}^m \beta_g \mathbb{E}_{P_g}[\ell(\theta; (X, Y))],$$

that involves continuous variables only, where $\Delta_{m-1} = \{\beta : \beta_g \geq 0, \sum_{g=1}^m \beta_g = 1\}$ is the $(m-1)$ -simplex.

Group DRO can be understood as an instance of standard DRO (Ben-Tal et al., 2013; Duchi et al., 2021) with a specific ambiguity set. Note that the standard DRO formulation is given as

$$\inf_{\theta \in \Theta} \sup_{Q \in \mathcal{Q}} \mathbb{E}_Q[\ell(\theta; (X, Y))], \quad (3)$$

where $\mathcal{Q}$ is a class of distributions, commonly referred to as the *ambiguity (or uncertainty) set*. The Group DRO formulation (2) is a specific case of DRO (3), with

$$\mathcal{Q} := \left\{ \sum_{g=1}^m \beta_g P_g : \beta \in \Delta_{m-1} \right\}. \quad (4)$$

Note that in frequently used DRO frameworks, the ambiguity set $\mathcal{Q}$ is often defined as a small neighborhood with respect to a standard (pseudo-)metric, such as the Wasserstein distance (Mohajerin Esfahani & Kuhn, 2018; Gao et al., 2024) and f -divergence (Namkoong & Duchi, 2016; Mehta et al., 2024).

4 PROPOSED METHOD

4.1 HIERARCHICAL AMBIGUITY SETS

In this section, we propose a hierarchical extension of Group DRO to be robust to distribution shifts at multiple levels. The proposed method is devised to capture both inter-group and intra-group uncertainties in modeling the distributional shifts.

High-level Formulation. As in Group DRO, we model the training distribution as a mixture of the form $P = \sum_{g=1}^m \alpha_g P_g$. To model the distributional uncertainty, we consider the DRO formulation (3) with a *hierarchical ambiguity set* $\mathcal{Q}$, defined as

$$\mathcal{Q} = \left\{ \sum_{g=1}^m \beta_g Q_g : \beta \in \Delta_{m-1}, d_1(\beta, \alpha) \leq \rho, d_2(Q_g, P_g) \leq \epsilon_g \quad \forall g \right\}, \quad (5)$$

where d_1 and d_2 are suitable metrics on Δ_{m-1} and the class of distributions for (X, Y) , respectively, and $\rho, \epsilon_g > 0$ are radii that determine the size of the ambiguity set.

The ambiguity set $\mathcal{Q}$ has a two-level hierarchy. The first level is controlled by the mixing proportion β . It accounts for uncertainty in the proportion of each subpopulation or group. Such uncertainty can arise, for example, if certain minority groups appear more frequently in evaluation settings than in the training set, thereby increasing their probability of occurrence and potentially amplifying spurious correlations if not properly addressed (Sugiyama & Storkey, 2006). At the second level, the distributional shift in each group is considered to capture within-group variability.

By jointly accounting for changes in the group proportion α and the conditional distributions $\{P_g\}_{g=1}^m$, the proposed framework provides two levels of robustness: inter-group generalization and resilience to intra-group variability. This dual modeling of real-world uncertainties enables the proposed method to address a broader range of distributional shifts compared to Group DRO (2) alone or standard DRO (3), which uses a standard (pseudo-)metric neighborhood as its ambiguity set.

Relationship to Group DRO and Standard DRO. The ambiguity set (4) used in Group DRO is a special case of the proposed ambiguity set (5). In particular, (4) can be obtained by setting $\rho = \infty$ and $\epsilon_g = 0$.

While standard DRO, which uses a standard metric neighborhood as an ambiguity set, can also be understood as a special case of the proposed method, the philosophy of the proposed ambiguity set differs from that of the standard ones. In standard DRO, the ambiguity set is taken as a small neighborhood with respect to a standard (pseudo-)metric. In contrast, we allow a large value for ρ , the radius that determines robustness to group proportions. Hence, distributions that are far from P with respect to standard metrics can also belong to the ambiguity set (5).

Detailed Formulation. The choice of d_1 is not critical because most reasonable metrics on Δ_{m-1} lead to similar behavior in our formulation. The constraint $d_1(\beta, \alpha) \leq \rho$ controls how much the test-time group proportions β are allowed to deviate from the training mixture α . When prior knowledge suggests that the test-time group proportions will not depart substantially from those observed in training (for example, they are expected to remain within a small neighborhood of α), one may choose a finite ρ to restrict attention to such mixtures. In the spurious-correlation setting we study, however, it is standard to assume no reliable prior on test-time group proportions and to evaluate robustness via worst-group accuracy, which can be viewed as examining performance under the extreme mixture that concentrates all probability mass on the worst-performing group. For this reason, in the remainder of this paper we instantiate our framework with $\rho = \infty$.

For d_2 , we use a Wasserstein distance. Other options such as f -divergences (e.g., KL divergence or total variation distance) are in principle also possible, but they induce ambiguity sets with a structural limitation in our setting. Ambiguity sets defined via an f -divergence require candidate distributions to be absolutely continuous with respect to the empirical distribution; since the empirical distribution has finite support, such sets can only reweight the observed training samples. In contrast, Wasserstein distances do not impose absolute continuity and therefore allow support shifts,

enabling the ambiguity set to include meaningful variations that may not appear in the training data. This choice is thus better aligned with the intra-group distributional shifts we aim to capture.

Among Wasserstein distances, we consider the infinite-order Wasserstein distance for computational convenience. Recall the definition of the Wasserstein distance of order $p \in [1, \infty)$:

$$W_p(Q, P) = \inf_{\gamma} \left\{ \left(\int c((x, y), (x', y'))^p d\gamma \right)^{\frac{1}{p}} \right\},$$

where the infimum is taken over every coupling γ of Q and P , and $c(\cdot, \cdot)$ is a cost function. The infinite-order Wasserstein distance is defined as $W_{\infty}(Q, P) = \sup_{p \geq 1} W_p(Q, P)$, with a variational representation

$$W_{\infty}(Q, P) = \inf \left\{ \epsilon > 0 : \begin{array}{l} Q(A) \leq P(A^{\epsilon}) \\ \text{for every Borel set } A \end{array} \right\}, \quad (6)$$

where A^{ϵ} denotes the ϵ -enlargement of A ; see Givens & Shortt (1984).

The cost function is defined in a latent semantic space, which is more effective than defining it in the space of raw data (Zeiler & Fergus, 2014; Krizhevsky et al., 2017). Specifically, we employ a deep neural network f^{θ} of depth L , defined as

$$f^{\theta}(x) = f_L^{\theta}(f_{L-1}^{\theta}(\dots f_1^{\theta}(x))),$$

and take the output of the $(L-1)$ -th layer (before the final fully connected layer) as the semantic representation:

$$z(x) := f_{L-1}^{\theta}(f_{L-2}^{\theta}(\dots f_1^{\theta}(x))). \quad (7)$$

We then define the cost function $c(\cdot, \cdot)$ as

$$c((x, y), (x', y')) = \begin{cases} \|z(x) - z(x')\|, & \text{if } y = y', \\ \infty, & \text{otherwise.} \end{cases}$$

Note that under our definition, $W_p(P, Q) = \infty$ if the marginals P and Q of Y differ. In all our applications, the group indicator G is defined as a pair (Y, A) ; hence, this definition does not cause any issues.

Proposed Hierarchical DRO Formulation. To sum up, the proposed hierarchical DRO can be written in the standard form (3) with the ambiguity set

$$\mathcal{Q} = \left\{ \sum_{g=1}^m \beta_g Q_g : \begin{array}{l} \beta \in \Delta_{m-1}, \\ W_{\infty}(Q_g, P_g) \leq \epsilon_g \quad \forall g \end{array} \right\}. \quad (8)$$

The flexibility in $\beta \in \Delta_{m-1}$ allows the group proportion to differ from α , which enables adaptation to new or changing subpopulation frequencies without introducing entirely new groups. With the constraint $W_{\infty}(Q_g, P_g) \leq \epsilon_g$, we accommodate plausible instance-level shifts within each group. Leveraging the semantic cost $c(\cdot, \cdot)$ allows the model to capture meaningful perturbations in high-dimensional feature spaces without conflating different class labels.

4.2 ALGORITHM

In this subsection, we provide an algorithm to solve the proposed DRO with the ambiguity set (8). We set $\epsilon_g = \epsilon / \sqrt{n_g}$, where n_g is the size of group g in the training data, and ϵ is a tunable hyperparameter that controls the degree of robustness to within-group distributional shifts. Intuitively, the fewer samples a group has, the more cautiously its potential distributional variations must be accounted for. The choice of ϵ is a common challenge across robust learning methods, particularly DRO-based approaches. While our framework remains effective across a broad range of values of ϵ , we propose a simple yet effective heuristic for its selection to further enhance the practicality of our framework. Further details on the selection of ϵ are provided in Appendix G.

Due to the hierarchical structure of the ambiguity set in (8), solving the resulting DRO problem is not straightforward. We therefore begin by reformulating the hierarchical DRO into a tractable optimization problem. We formally state the resulting formulation in the following theorem. The proof is provided in Appendix A.

Theorem 4.1. *Let $\mathcal{Q}$ be the ambiguity set defined in (8). Then, the corresponding distributionally robust optimization problem*

$$\inf_{\theta \in \Theta} \sup_{Q \in \mathcal{Q}} \mathbb{E}_Q[\ell(\theta; (X, Y))]$$

is upper-bounded by the following surrogate objective:

$$\inf_{\theta \in \Theta} \sup_{\beta \in \Delta_{m-1}} \sum_{g=1}^m \beta_g \mathbb{E}_{P_g} \left[\sup_{z': \|z' - z(X)\| \leq \epsilon_g} \mathcal{L}(f_L^\theta(z'), Y) \right]. \quad (9)$$

Intuitively, Theorem 4.1 shows that the worst-case risk over our hierarchical ambiguity set can be conservatively over-approximated by an adversarial perturbation problem in the latent space, where the inner maximization is weighted by the worst-case group proportions β .

Remark (Tightness of the relaxation). The upper bound in (9) relies on relaxing the inner maximization from

$$\mathbb{E}_{P_g} \left[\sup_{x: \|z(x) - z(X)\| \leq \epsilon_g} \mathcal{L}(f_L^\theta(z(x)), Y) \right] \leq \mathbb{E}_{P_g} \left[\sup_{z': \|z' - z(X)\| \leq \epsilon_g} \mathcal{L}(f_L^\theta(z'), Y) \right]$$

This relaxation can be strict in general when the feature map $z(\cdot)$ is not surjective onto the ϵ_g -ball. However, we expect this effect to be negligible in practice for embedding distributions arising in modern deep networks. It is commonly assumed that the pushforward distribution of X under $z(\cdot)$ possesses a positive Lebesgue density on, or in a neighborhood of, its effective support.

Although raw images may concentrate near a low-dimensional manifold in input space and thus typically do not admit a Lebesgue density, the situation is different after they pass through several nonlinear layers. Empirically and conceptually, internal representations produced by deep networks behave as “thickened” manifolds and are typically modeled as distributions with nonzero volume in the ambient feature space.

Formally proving the existence of such a density for arbitrary deep embeddings is challenging. Nevertheless, under this standard modeling assumption, the relaxation above becomes nearly tight. The overlap between the ϵ_g -ball and the image of $z(\cdot)$ is substantial, making the resulting upper bound in (9) a close approximation of the original hierarchical objective.

We therefore minimize the surrogate objective (9) via a coordinate-wise procedure, as detailed next.

Proposed Iterative Training Procedure. For a given θ , let z'_i denote the maximizer of the map

$$z' \mapsto \mathcal{L}(f_L^\theta(z'), y_i)$$

over the set $\{z' : \|z' - z(x_i)\| \leq \epsilon_{g_i}\}$. To solve the optimization problem (9), we iteratively update β , θ and semantic variables z'_i coordinate-wise as below. A pseudo-code for a minibatch size of 1 is provided in Algorithm 1.

1. *Update of z' .* For given θ , z'_i can be approximated by one-step projected gradient ascent, ensuring that $\|z' - z(x_i)\| \leq \epsilon_{g_i}$. (Lines 6–8)
2. *Update of β .* For given θ and z'_i , β can be computed using exponentiated gradient ascent, a variant of mirror descent with negative Shannon entropy (Sagawa et al., 2020). (Lines 10–12)
3. *Update of θ .* For a given β and z'_i , we update θ using stochastic gradient descent. (Line 13)

A convergence guarantee under convexity assumptions is established in Appendix B, showing that the algorithm achieves an $O(1/\sqrt{T})$ convergence rate.

5 EXPERIMENTS

5.1 DATASET

We conduct experiments on three widely used benchmark datasets, CMNIST, Waterbirds, and CelebA, each exhibiting known spurious correlations between the label and an irrelevant attribute. All datasets include a minority group that is underrepresented, rendering them susceptible to distributional shifts.

Algorithm 1 DRO with a Hierarchical Ambiguity Set

```

1: Input: Step sizes  $\eta_\beta, \eta_\theta, \eta_z$ ; initial parameters  $\theta^{(0)}, \beta^{(0)}$ ; number of iterations  $T$ 
2: for  $t = 1$  to  $T$  do
3:   Sample  $g \sim \text{Uniform}(1, \dots, m)$ 
4:   Sample  $(x, y) \sim P_g$ 
5:   Initialize  $z' \leftarrow z(x)$ 
6:    $z' \leftarrow z' + \eta_z \nabla_{z'} \mathcal{L}(f_L^{\theta^{(t-1)}}(z'), y)$ 
7:   if  $\|z' - z(x)\| > \epsilon_g$  then
8:      $z' \leftarrow \text{Proj}_{\|z' - z(x)\| \leq \epsilon_g}(z')$ 
9:   end if
10:  Update  $\beta' \leftarrow \beta^{(t-1)}$ 
11:  Update  $\beta'_g \leftarrow \beta'_g \exp\left(\eta_\beta \mathcal{L}(f_L^{\theta^{(t-1)}}(z'), y)\right)$ 
12:  Normalize  $\beta^{(t)} \leftarrow \beta' / \sum \beta'_{g'}$ 
13:  Update  $\theta^{(t)} \leftarrow \theta^{(t-1)} - \eta_\theta \beta_g^{(t)} \nabla_\theta \mathcal{L}(f_L^{\theta^{(t-1)}}(z'), y)$ 
14: end for

```

Original Datasets.

- **CMNIST** (Arjovsky et al., 2019): A colored variant of MNIST, split into four groups based on digit label (*digits 0–4 as label 0*, and *digits 5–9 as label 1*) and color (*red vs. green*). The color is spuriously correlated with the digit label in the training set.
- **Waterbirds** (Sagawa et al., 2020): Created by combining bird images from CUB (Wah et al., 2011) with backgrounds from Places (Zhou et al., 2017), yielding four groups based on (*bird type, background*). The minority group (*waterbird, land background*) typically has few samples.
- **CelebA** (Liu et al., 2015): A facial attribute dataset used here for classifying *blond vs. non-blond* hair, where *gender* acts as a spurious attribute. The minority group (*blond hair, male*) is significantly underrepresented.

Modified Datasets with Minority Group Shifts. To rigorously test our approach under more realistic distribution shifts, we construct modified versions of the above datasets by inducing intra-group shifts specifically in each minority group:

- **Shifted CMNIST:** Rotate all images in the minority group (*label 1, red*) by 90° at test time, while keeping them unrotated at training time.
- **Shifted Waterbirds:** Restrict the training set’s minority group (*waterbird, land background*) to only *waterfowls*, and the test set’s minority group to only *seabirds*.
- **Shifted CelebA:** For minority group (*blond hair, male*), include only *no-glasses* images in training and only *with-glasses* images at test time.

These modifications reflect real-world scenarios where underrepresented groups not only appear more frequently but also exhibit subtle changes. Further details and illustrative examples are provided in Appendix D.

5.2 BASELINES

We compare our method to several representative baselines: ERM, Group DRO (Sagawa et al., 2020), JTT (Liu et al., 2021), CnC (Zhang et al., 2022), SSA (Nam et al., 2022), LISA (Yao et al., 2022), DFR (Kirichenko et al., 2023), PDE (Deng et al., 2024), and GIC (Han & Zou, 2024). These methods range from direct robust learning (e.g., Group DRO) to two-step pipelines that first infer group membership and then apply robust training (e.g., SSA, GIC). Detailed descriptions are provided in Appendix E.

For our newly constructed datasets incorporating minority-group distribution shifts, we conducted experiments focusing on Group DRO, LISA, DFR, and PDE. Unlike methods that infer group labels and then rely on a separate robust training step, these four baselines—like our proposed approach—directly utilize known group information. This distinction provides a more consistent and fair comparison in scenarios where explicit group labels are available.

Table 1: Worst-group and average accuracy on CMNIST, Waterbirds, and CelebA under shifted distributions. All results are averaged over three runs with different random seeds. Boldface indicates the best performance, while underlined numbers denote the second-best.

Method	Group label	Shifted CMNIST		Shifted Waterbirds		Shifted CelebA	
		Worst Acc	Average Acc	Worst Acc	Average Acc	Worst Acc	Average Acc
GroupDRO	✓	65.9±8.2	74.0±0.7	91.7±0.3	94.9±0.1	59.8±3.2	92.4±0.3
LISA	✓	42.9±10.0	59.8±4.5	79.1±1.8	94.2±0.3	<u>60.6</u> ±1.1	92.1±0.2
DFR ^{tr}	✓	28.0±4.9	47.8±1.9	89.2±1.5	96.3±0.4	50.3±3.5	90.5±0.4
PDE	✓	65.3±11.1	71.3±6.2	84.4±4.6	92.0±0.6	56.3±11.2	91.6±0.4
Ours	✓	71.8 ±2.8	75.0±0.4	93.7 ±0.2	94.6±0.1	72.1 ±2.0	91.3±0.1

Table 2: Worst-group and average accuracy on CMNIST, Waterbirds, and CelebA under their original (un-shifted) distributions.

Method	Group label	CMNIST		Waterbirds		CelebA	
		Worst Acc	Average Acc	Worst Acc	Average Acc	Worst Acc	Average Acc
ERM	×	3.4±0.9	12.9±0.8	62.6±0.3	97.3±1.0	47.7±2.1	94.9±0.3
JTT	×	67.3±5.1	76.4±3.3	83.8±1.2	89.3±0.7	81.5±1.7	88.1±0.3
CnC	×	—	—	88.5±0.3	90.9±0.1	88.8±0.9	89.9±0.5
GIC	×	72.2±0.5	73.2±0.2	86.3±0.1	89.6±1.3	89.4±0.2	91.9±0.1
SSA	×	71.1±0.4	75.0±0.3	89.0±0.6	92.2±0.9	89.8±1.3	92.8±0.1
GroupDRO	✓	73.1±0.3	74.8±0.2	<u>90.6</u> ±0.2	92.7±0.1	89.3±1.3	92.6±0.3
LISA	✓	73.3±0.2	74.0±0.1	89.2±0.6	91.8±0.3	89.3±1.1	92.4±0.4
DFR ^{tr}	✓	59.8±0.4	62.1±0.2	90.2±0.8	97.0±0.3	80.7±2.4	90.6±0.7
PDE	✓	72.6±0.7	73.0±0.4	90.3±0.3	92.4±0.8	91.0 ±0.4	92.0±0.6
Ours	✓	73.6 ±0.3	75.1±0.5	90.8 ±0.2	92.6±0.2	<u>90.4</u> ±0.3	92.7±0.0

5.3 EVALUATION

Metrics. We consider two metrics: *worst-group accuracy* and *average accuracy*. The worst-group accuracy is obtained by evaluating accuracy on each group and taking the minimum across all groups, providing insight into how a method performs if the test distribution is heavily skewed toward the most challenging subgroup. Meanwhile, the average accuracy is computed as the weighted average of group accuracy, where the weights are proportional to the group sizes in the training data, reflecting overall performance but offering less visibility into group-specific disparities.

Model Selection. Following Sagawa et al. (2020) and related methods, we select hyperparameters and stopping criteria based on the highest worst-group validation accuracy. In particular, for scenarios involving minority group shifts, we adopt the data-driven tuning procedure from Appendix G to determine the perturbation parameter ϵ .

5.4 RESULTS

Performance on Shifted Distributions. Under shifted distributions (Table 1), our method demonstrates clear superiority in worst-group accuracy across all three benchmarks (CMNIST, Waterbirds, and CelebA). On CMNIST, LISA and DFR degrade substantially, highlighting their vulnerability to intra-group shifts. By contrast, our framework maintains high worst-group accuracy.

For Waterbirds, which involves a moderate shift in species composition within the minority group, most baselines experience notable drops in worst-group accuracy. In contrast, our approach maintains robust worst-group accuracy, indicating its capacity to adapt to intra-group variability. Interestingly, both Group DRO and our method report higher worst-group accuracy in the shifted case than in the original Waterbirds dataset (Table 2); however, this discrepancy arises from a known mislabeling issue (Asgari et al., 2022), where three bird species labeled as “waterbird” should actually be “landbird.” To verify this, we correct the mislabeled samples in the original dataset and report results

in Table 3: under the corrected labels, both Group DRO and our method exhibit the expected pattern, performing better in the unshifted setting than under the minority-group shift. Notably, our approach outperforms Group DRO in both scenarios, confirming its robustness even after label corrections.

On the more challenging CelebA benchmark, our advantage grows more pronounced. While PDE shows slightly higher worst-group accuracy on the original dataset (Table 2), its performance drops sharply (by about 34.7%) when the minority-group distribution is shifted. These observations underscore the importance of modeling both inter-group and intra-group uncertainties—especially given that minority groups in Waterbirds and CelebA constitute only about 1% of the data and thus are more susceptible to distributional changes. Furthermore, our results highlight that relying on pre-defined test splits with uniformly distributed attributes may offer an overly optimistic view of real-world robustness.

Table 3: Worst-group and average accuracy on Waterbirds under minority group shifted distributions and Corrected Waterbirds on the original dataset with corrected labels.

Method	Shifted Waterbirds		Corrected Waterbirds	
	Worst Acc	Avg Acc	Worst Acc	Avg Acc
Group DRO	91.7 $\pm$ 0.3	94.9 $\pm$ 0.1	94.1 $\pm$ 0.6	94.7 $\pm$ 0.0
Ours	93.7$\pm$0.2	94.6 $\pm$ 0.1	95.1$\pm$0.4	96.3 $\pm$ 0.0

Performance on Original Distributions. On the original (unshifted) versions of CMNIST, Waterbirds, and CelebA (Table 2), our method consistently achieves top-tier worst-group accuracy. It secures the highest or near-highest scores across all three benchmarks, confirming that the proposed framework not only excels under distributional shifts but also remains effective when intra-group distributions are stable. Notably, even in these unshifted settings, methods such as Group DRO rely on the strong assumption that each group’s training distribution remains valid at test time. As our results show, explicitly modeling distributional uncertainty within minority groups can yield more reliable robustness, highlighting the limitations of approaches that treat group distributions as fixed. By addressing potential discrepancies at both the inter-group and intra-group levels, our framework provides a stronger foundation for real-world applications.

6 CONCLUSION

Summary. We introduce a novel distributionally robust optimization framework with a hierarchical ambiguity set that explicitly models both inter-group and intra-group distribution shifts—a previously overlooked but practically crucial scenario for underrepresented subpopulations. We reformulate the resulting DRO into a tractable and computationally efficient optimization problem, making it feasible for practical use. Our analysis shows that even modest changes in how minority group samples are partitioned between training and testing can severely degrade the performance of existing robust methods. In contrast, our approach consistently maintains strong performance on both standard and modified benchmarks.

Outlook. While our experiments focus on image datasets with explicit feature labels that make intra-group shifts observable, extending our framework to domains such as text—where such labels are difficult to obtain—represents an interesting avenue for future work. More broadly, we hope this work motivates future efforts to systematically uncover and address overlooked failure modes in robust learning.

ACKNOWLEDGMENTS

This work was supported by the National Research Foundation of Korea (NRF) grant funded by the Korea government (MSIT) (No. RS-2023-00240861), a Korea Institute for Advancement of Technology (KIAT) grant funded by the Korea Government (MOTIE) (RS-2024-00409092, 2024 HRD Program for Industrial Innovation), and Institute of Information & Communications Technology Planning & Evaluation(IITP)-Global Data-X Leader HRD program grant funded by the Korea government (MSIT) (IITP-2024-RS-2024-00441244).

REFERENCES

- Faruk Ahmed, Yoshua Bengio, Harm Van Seijen, and Aaron Courville. Systematic generalisation with group invariant predictions. In *Proc. International Conference on Learning Representations*, 2021.
- M. Arjovsky, L. Bottou, I. Gulrajani, and D. Lopez-Paz. Invariant risk minimization. *arXiv preprint arXiv:1907.02893*, 2019.
- Saeid Asgari, Aliasghar Khani, Fereshte Khani, Ali Gholami, Linh Tran, Ali Mahdavi Amiri, and Ghassan Hamarneh. Masktune: Mitigating spurious correlations by forcing to explore. In *Proc. Advances in Neural Information Processing Systems*, 2022.
- Marcus A Badgeley, John R Zech, Luke Oakden-Rayner, Benjamin S Glicksberg, Manway Liu, William Gale, Michael V McConnell, Bethany Percha, Thomas M Snyder, and Joel T Dudley. Deep learning predicts hip fracture using confounding patient and healthcare variables. *NPJ Digital Medicine*, 2(1):31, 2019.
- Xingjian Bai, Guangyi He, Yifan Jiang, and Jan Obloj. Wasserstein distributional robustness of neural networks. In *Proc. Advances in Neural Information Processing Systems*, 2024.
- A. Ben-Tal, D. Den Hertog, A. De Waegenaere, B. Melenberg, and G. Rennen. Robust solutions of optimization problems affected by uncertain probabilities. *Management Science*, 2013.
- Aharon Ben-Tal and Arkadi Nemirovski. Robust optimization-methodology and applications. *Mathematical Programming*, 92:453–480, 2002.
- Jose Blanchet, Lin Chen, and Xun Yu Zhou. Distributionally robust mean-variance portfolio selection with wasserstein distances. *Management Science*, 68(9):6382–6410, 2022.
- Joy Buolamwini and Timnit Gebru. Gender shades: Intersectional accuracy disparities in commercial gender classification. In *Proc. Conference on Fairness, Accountability and Transparency*, pp. 77–91, 2018.
- Tianle Cai, Ruiqi Gao, Jason Lee, and Qi Lei. A theory of label propagation for subpopulation shift. In *Proc. International Conference on Machine Learning*, pp. 1170–1182, 2021.
- Elliot Creager, Jörn-Henrik Jacobsen, and Richard Zemel. Environment inference for invariant learning. In *Proc. International Conference on Machine Learning*, pp. 2189–2200, 2021.
- Yihe Deng, Yu Yang, Baharan Mirzasoleiman, and Quanquan Gu. Robust learning with progressive data expansion against spurious correlation. In *Proc. Advances in Neural Information Processing Systems*, 2024.
- Luc Devroye, László Györfi, and Gábor Lugosi. *A Probabilistic Theory of Pattern Recognition*. Springer Science & Business Media, 2013.
- J. C. Duchi, P. W. Glynn, and H. Namkoong. Statistics of robust optimization: A generalized empirical likelihood approach. *Mathematics of Operations Research*, 2021.
- John C Duchi and Hongseok Namkoong. Learning models with uniform performance via distributionally robust optimization. *Annals of Statistics*, 49(3):1378–1406, 2021.
- Rui Gao, Xi Chen, and Anton J Kleywegt. Wasserstein distributionally robust optimization and variation regularization. *Operations Research*, 72(3):1177–1191, 2024.

- Soumya Suvra Ghosal and Yixuan Li. Distributionally robust optimization with probabilistic group. In *Proc. AAAI Conference on Artificial Intelligence*, pp. 11809–11817, 2023.
- Clark R Givens and Rae Michael Shortt. A class of wasserstein metrics for probability distributions. *The Michigan Mathematical Journal*, 31(2):231–240, 1984.
- Yujin Han and Difan Zou. Improving group robustness on spurious correlation requires preciser group inference. In *Proc. International Conference on Machine Learning*, pp. 17480–17504, 2024.
- Tatsunori Hashimoto, Megha Srivastava, Hongseok Namkoong, and Percy Liang. Fairness without demographics in repeated loss minimization. In *Proc. International Conference on Machine Learning*, pp. 1929–1938, 2018.
- Pavel Izmailov, Polina Kirichenko, Nate Gruver, and Andrew G Wilson. On feature learning in the presence of spurious correlations. In *Proc. Advances in Neural Information Processing Systems*, 2022.
- Saachi Jain, Kimia Hamidieh, Kristian Georgiev, Andrew Ilyas, Marzyeh Ghassemi, and Aleksander Madry. Improving subgroup robustness via data selection. In *Proc. Advances in Neural Information Processing Systems*, 2024.
- Nayeong Kim, Juwon Kang, Sungsoo Ahn, Jungseul Ok, and Suha Kwak. Improving robustness to multiple spurious correlations by multi-objective optimization. In *Proc. International Conference on Machine Learning*, pp. 24040–24058, 2024.
- Seonghwi Kim, Sung Ho Jo, Wooseok Ha, and Minwoo Chae. Distributionally robust classification for multi-source unsupervised domain adaptation. In *Proc. International Conference on Learning Representations*, 2026.
- Polina Kirichenko, Pavel Izmailov, and Andrew Gordon Wilson. Last layer re-training is sufficient for robustness to spurious correlations. In *Proc. International Conference on Learning Representations*, 2023.
- P. W. Koh, S. Sagawa, H. Marklund, S. M. Xie, M. Zhang, A. Balsubramani, W. Hu, M. Yasunaga, R. L. Phillips, I. Gao, et al. Wilds: A benchmark of in-the-wild distribution shifts. In *Proc. International Conference on Machine Learning*, pp. 5637–5664, 2021.
- A. Krizhevsky, I. Sutskever, and G. E. Hinton. Imagenet classification with deep convolutional neural networks. *Communications of the ACM*, 60(6):84–90, 2017.
- David Krueger, Ethan Caballero, Joern-Henrik Jacobsen, Amy Zhang, Jonathan Binas, Dinghui Zhang, Remi Le Priol, and Aaron Courville. Out-of-distribution generalization via risk extrapolation (rex). In *Proc. International Conference on Machine Learning*, pp. 5815–5826, 2021.
- Daniel Kuhn, Peyman Mohajerin Esfahani, Viet Anh Nguyen, and Soroosh Shafieezadeh-Abadeh. Wasserstein distributionally robust optimization: Theory and applications in machine learning. In *Operations Research & Management Science in the Age of Analytics*, pp. 130–166. Informs, 2019.
- Tyler LaBonte, Vidya Muthukumar, and Abhishek Kumar. Towards last-layer retraining for group robustness with fewer annotations. In *Proc. Advances in Neural Information Processing Systems*, 2024.
- Evan Z Liu, Behzad Haghighi, Annie S Chen, Aditi Raghunathan, Pang Wei Koh, Shiori Sagawa, Percy Liang, and Chelsea Finn. Just train twice: Improving group robustness without training group information. In *Proc. International Conference on Machine Learning*, pp. 6781–6792, 2021.
- Z. Liu, P. Luo, X. Wang, and X. Tang. Deep learning face attributes in the wild. In *Proc. International Conference on Computer Vision*, pp. 3730–3738, 2015.

- Ronak Mehta, Vincent Roulet, Krishna Pillutla, and Zaid Harchaoui. Distributionally robust optimization with bias and variance reduction. In *Proc. International Conference on Learning Representations*, 2024.
- 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.
- Varun Mulchandani and Jung-Eun Kim. Severing spurious correlations with data pruning. In *Proc. International Conference on Learning Representations*, 2025.
- Junhyun Nam, Hyuntak Cha, Sungsoo Ahn, Jaeho Lee, and Jinwoo Shin. Learning from failure: De-biasing classifier from biased classifier. In *Proc. Advances in Neural Information Processing Systems*, 2020.
- Junhyun Nam, Jaehyung Kim, Jaeho Lee, and Jinwoo Shin. Spread spurious attribute: Improving worst-group accuracy with spurious attribute estimation. In *Proc. International Conference on Learning Representations*, 2022.
- Hongseok Namkoong and John C Duchi. Stochastic gradient methods for distributionally robust optimization with f-divergences. In *Proc. Advances in Neural Information Processing Systems*, 2016.
- Arkadi Nemirovski, Anatoli Juditsky, Guanghui Lan, and Alexander Shapiro. Robust stochastic approximation approach to stochastic programming. *SIAM Journal on optimization*, 19(4):1574–1609, 2009.
- Ziad Obermeyer, Brian Powers, Christine Vogeli, and Sendhil Mullainathan. Dissecting racial bias in an algorithm used to manage the health of populations. *Science*, 366(6464):447–453, 2019.
- Bhargavi Paranjape, Pradeep Dasigi, Vivek Srikumar, Luke Zettlemoyer, and Hannaneh Hajishirzi. AGRO: Adversarial discovery of error-prone groups for robust optimization. In *Proc. International Conference on Learning Representations*, 2023.
- Vihari Piratla, Praneeth Netrapalli, and Sunita Sarawagi. Focus on the common good: Group distributional robustness follows. In *Proc. International Conference on Learning Representations*, 2022.
- Shikai Qiu, Andres Potapczynski, Pavel Izmailov, and Andrew Gordon Wilson. Simple and fast group robustness by automatic feature reweighting. In *Proc. International Conference on Machine Learning*, pp. 28448–28467, 2023.
- S. Sagawa, P. W. Koh, T. B. Hashimoto, and P. Liang. Distributionally robust neural networks for group shifts: On the importance of regularization for worst-case generalization. In *Proc. International Conference on Learning Representations*, 2020.
- Ramprasaath R Selvaraju, Michael Cogswell, Abhishek Das, Ramakrishna Vedantam, Devi Parikh, and Dhruv Batra. Grad-cam: Visual explanations from deep networks via gradient-based localization. In *Proc. International Conference on Computer Vision*, pp. 618–626, 2017.
- Seonguk Seo, Joon-Young Lee, and Bohyung Han. Unsupervised learning of debiased representations with pseudo-attributes. In *Proc. Conference on Computer Vision and Pattern Recognition*, pp. 16742–16751, 2022.
- Nimit Sohoni, Jared Dunnmon, Geoffrey Angus, Albert Gu, and Christopher Ré. No subclass left behind: Fine-grained robustness in coarse-grained classification problems. In *Proc. Advances in Neural Information Processing Systems*, 2020.
- M. Staib and S. Jegelka. Distributionally robust deep learning as a generalization of adversarial training. In *NeurIPS Workshop on Machine Learning and Computer Security*, 2017.
- M. Sugiyama and A. J. Storkey. Mixture regression for covariate shift. In *Proc. Advances in Neural Information Processing Systems*, 2006.

- Laurens Van der Maaten and Geoffrey Hinton. Visualizing data using t-SNE. *Journal of Machine Learning Research*, 9(86):2579–2605, 2008.
- Riccardo Volpi, Hongseok Namkoong, Ozan Sener, John C Duchi, Vittorio Murino, and Silvio Savarese. Generalizing to unseen domains via adversarial data augmentation. In *Proc. Advances in Neural Information Processing Systems*, 2018.
- C. Wah, S. Branson, P. Welinder, P. Perona, and S. Belongie. The caltech-ucsd birds-200-2011 dataset. Technical Report CNS-TR-2011-001, California Institute of Technology, 2011.
- Shirley Wu, Mert Yuksekgonul, Linjun Zhang, and James Zou. Discover and cure: Concept-aware mitigation of spurious correlation. In *Proc. International Conference on Machine Learning*, pp. 37765–37786, 2023.
- Yuzhe Yang, Haoran Zhang, Dina Katabi, and Marzyeh Ghassemi. Change is hard: A closer look at subpopulation shift. In *Proc. International Conference on Machine Learning*, pp. 39584–39622, 2023.
- Huaxiu Yao, Yu Wang, Sai Li, Linjun Zhang, Weixin Liang, James Zou, and Chelsea Finn. Improving out-of-distribution robustness via selective augmentation. In *Proc. International Conference on Machine Learning*, pp. 25407–25437, 2022.
- John R Zech, Marcus A Badgeley, Manway Liu, Anthony B Costa, Joseph J Titano, and Eric Karl Oermann. Variable generalization performance of a deep learning model to detect pneumonia in chest radiographs: A cross-sectional study. *PLoS Medicine*, 15(11):e1002683, 2018.
- M. D. Zeiler and R. Fergus. Visualizing and understanding convolutional networks. In *Proc. European Conference on Computer Vision*, pp. 818–833, 2014.
- Michael Zhang, Nimit S Sohoni, Hongyang R Zhang, Chelsea Finn, and Christopher Re. Correct-N-Contrast: A contrastive approach for improving robustness to spurious correlations. In *Proc. International Conference on Machine Learning*, pp. 26484–26516, 2022.
- Yang Zhang, Philip David, and Boqing Gong. Curriculum domain adaptation for semantic segmentation of urban scenes. In *Proc. International Conference on Computer Vision*, pp. 2020–2030, 2017.
- B. Zhou, A. Lapedriza, A. Khosla, A. Oliva, and A. Torralba. Places: A 10 million image database for scene recognition. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, pp. 1452–1464, 2017.
- Chunting Zhou, Xuezhe Ma, Paul Michel, and Graham Neubig. Examining and combating spurious features under distribution shift. In *Proc. International Conference on Machine Learning*, pp. 12857–12867, 2021.
- Reihaneh Zohrabi, Hosein Hasani, Mahdiah Soleymani Baghshah, Anna Rohrbach, Marcus Rohrbach, and Mohammad Hossein Rohban. Spurious-aware prototype refinement for reliable out-of-distribution detection. *arXiv preprint arXiv:2506.23881*, 2025.

LLM USAGE

We used large language models solely to aid in polishing the writing of this paper.

A PROOF OF THEOREM 4.1

Proof. We begin with a lemma adapted from Staib & Jegelka (2017), with minor adjustments to match our framework. This lemma provides an equivalent form for the inner supremum problem of DRO with a W_∞ -neighborhood, which is closely related to the representation (6) of W_∞ .

Lemma A.0.1. (Staib & Jegelka, 2017, Proposition 3.1) *Let θ be fixed model parameters, and let $c(\cdot, \cdot)$ be a metric on the input space $\mathcal{X}$. For any distribution P on $\mathcal{X} \times \mathcal{Y}$ and for any $\epsilon \geq 0$,*

$$\mathbb{E}_P \left[\sup_{(x,y) \in B_\epsilon(X,Y)} \ell(\theta; (x,y)) \right] = \sup_{W_\infty(Q,P) \leq \epsilon} \mathbb{E}_Q[\ell(\theta; (X,Y))].$$

where $B_\epsilon(x,y) = \{(x',y') : c((x,y), (x',y')) \leq \epsilon\}$.

With the ambiguity set (8), the DRO (3) is equivalent to

$$\inf_{\theta \in \Theta} \left\{ \sup_{\beta \in \Delta_{m-1}} \sup_{\substack{W_\infty(Q_g, P_g) \leq \epsilon_g \\ g=1, \dots, m}} \mathbb{E}_Q[\ell(\theta; (X,Y))] \right\}, \quad (10)$$

where $Q = \sum_{g=1}^m \beta_g Q_g$.

For a fixed θ , the double supremum in (10) can equivalently be written as

$$\begin{aligned} & \sup_{\beta \in \Delta_{m-1}} \sup_{\substack{W_\infty(Q_g, P_g) \leq \epsilon_g \\ g=1, \dots, m}} \sum_{g=1}^m \beta_g \mathbb{E}_{Q_g}[\ell(\theta; (X,Y))] \\ &= \sup_{\beta \in \Delta_{m-1}} \sum_{g=1}^m \beta_g \sup_{W_\infty(Q_g, P_g) \leq \epsilon_g} \mathbb{E}_{Q_g}[\ell(\theta; (X,Y))]. \end{aligned}$$

By applying Lemma A.0.1, one can upper-bound the inner supremum in the previous display as

$$\begin{aligned} \mathbb{E}_{P_g} \left[\sup_{x: \|z(x) - z(X)\| \leq \epsilon_g} \ell(\theta; (x,Y)) \right] &= \mathbb{E}_{P_g} \left[\sup_{x: \|z(x) - z(X)\| \leq \epsilon_g} \mathcal{L}(f_L^\theta(z(x)), Y) \right] \\ &\leq \mathbb{E}_{P_g} \left[\sup_{z': \|z' - z(X)\| \leq \epsilon_g} \mathcal{L}(f_L^\theta(z'), Y) \right], \end{aligned}$$

where $x \mapsto z(x)$ denotes the feature map defined in (7). Thus, the original optimization problem is upper-bounded by

$$\inf_{\theta \in \Theta} \sup_{\beta \in \Delta_{m-1}} \sum_{g=1}^m \beta_g \mathbb{E}_{P_g} \left[\sup_{z': \|z' - z(X)\| \leq \epsilon_g} \mathcal{L}(f_L^\theta(z'), Y) \right],$$

and completes the proof. $\square$

B CONVERGENCE ANALYSIS OF ALGORITHM 1

We analyze convergence via ε_T of the average iterate $\bar{\theta}^{(1:T)}$:

$$\varepsilon_T = \max_{\beta \in \Delta_{m-1}} L(\bar{\theta}^{(1:T)}, \beta) - \min_{\theta \in \Theta} \max_{\beta \in \Delta_{m-1}} L(\theta, \beta),$$

where $L(\theta, \beta) := \sum_{g=1}^m \beta_g \mathbb{E}_{P_g} \left[\sup_{z': \|z' - z(X)\| \leq \epsilon_g} \mathcal{L}(f_L^\theta(z'), Y) \right]$. In the convex setting, our method achieves $O(1/\sqrt{T})$.

Proposition B.1 (Convergence of Algorithm 1). *Suppose $\mathcal{L}(f_L^\theta(z), y)$ is non-negative, convex in θ , B_∇ -Lipschitz in θ , and bounded by B_ℓ for all (x, y) in $\mathcal{X} \times \mathcal{Y}$. In addition, let $\|\theta\|_2 \leq B_\Theta$ for all θ in some convex set $\Theta \subset \mathbb{R}^d$, and assume the feature map $z(x)$ is fixed w.r.t. θ . Then, the average iterate of Algorithm 1 achieves an expected error at the rate*

$$\mathbb{E}[\varepsilon_T] \leq 2m \sqrt{\frac{10(B_\Theta^2 B_\nabla^2 + B_\ell^2 \log m)}{T}}.$$

Proof. Each iteration samples $G \sim \text{Unif}\{1, \dots, m\}$ and $(X, Y) \sim P_G$. The resulting joint sample $\xi = (X, Y, G)$ is drawn i.i.d. from the mixture distribution $q := \frac{1}{m} \sum_{g=1}^m P_g$.

For each group $g \in \{1, \dots, m\}$, define the stochastic loss function

$$F_g(\theta; \xi) := m \cdot \mathbf{1}[G = g] \cdot \sup_{\|z' - z(X)\| \leq \epsilon_g} \mathcal{L}(f_L^\theta(z'), Y),$$

and let

$$f_g(\theta) := \mathbb{E}_{P_g} \left[\sup_{\|z' - z(X)\| \leq \epsilon_g} \mathcal{L}(f_L^\theta(z'), Y) \right].$$

The total objective is then $L(\theta, \beta) = \sum_{g=1}^m \beta_g f_g(\theta)$.

We now verify the conditions required to apply the standard online mirror descent (OMD) regret bound (Nemirovski et al., 2009):

- (A) Convexity. For each g , the inner function $\mathcal{L}(f_L^\theta(z'), Y)$ is convex and non-negative in θ , and the supremum preserves convexity via Danskin's theorem. Thus, $f_g(\theta)$ is convex.
- (B) Expectation form. We have

$$\mathbb{E}_{\xi \sim q} [F_g(\theta; \xi)] = \frac{1}{m} \sum_{g'=1}^m \mathbb{E}_{(X, Y) \sim P_{g'}} \left[m \cdot \mathbf{1}[g' = g] \cdot \sup_{\|z' - z(X)\| \leq \epsilon_g} \mathcal{L}(f_L^\theta(z'), Y) \right] = f_g(\theta).$$

- (C) Unbiased subgradients. By Danskin's theorem, the mapping $\theta \mapsto \sup_{z'} \mathcal{L}(f_L^\theta(z'), Y)$ is subdifferentiable. Hence, $\nabla_\theta F_g(\theta; \xi)$ is an unbiased subgradient:

$$\mathbb{E}_{\xi \sim q} [\nabla_\theta F_g(\theta; \xi)] = \nabla_\theta f_g(\theta).$$

With the conditions (A)–(C) established, and using the boundedness assumptions:

$$\|\theta\|_2 \leq B_\Theta, \quad \|\nabla_\theta \mathcal{L}\| \leq B_\nabla, \quad \mathcal{L} \leq B_\ell,$$

the standard OMD regret bound (Nemirovski et al., 2009; Sagawa et al., 2020) yields

$$\mathbb{E}[\varepsilon_T] \leq 2m \sqrt{\frac{10(B_\Theta^2 B_\nabla^2 + B_\ell^2 \log m)}{T}},$$

completing the proof. $\square$

The convergence guarantee in Proposition B.1 is derived under the assumption that the feature map $z(\cdot)$ is fixed with respect to θ . Under this assumption, the resulting optimization problem becomes convex in θ , which allows us to establish the stated $O(1/\sqrt{T})$ convergence rate. Accordingly, Proposition B.1 should be interpreted as characterizing the behavior of Algorithm 1 in the idealized regime where the representation is fixed, rather than as a guarantee for the full nonconvex joint optimization of all network parameters.

C INTERPRETING LATENT PERTURBATION REGULARIZATION

To clarify the intuition behind our latent perturbation framework, we employ a first-order Taylor expansion of the loss function. This approximation shows that the innermost supremum in our optimization problem can be interpreted as the original loss $\mathcal{L}(f^\theta(x), y)$ plus an additional regularization term involving the dual norm of the gradient with respect to the latent representation. Specifically,

$$\begin{aligned} \sup_{\|z' - z(x)\| \leq \epsilon} \mathcal{L}(f_L^\theta(z'), y) &= \sup_{\|z' - z(x)\| \leq \epsilon} \mathcal{L}(f_L^\theta(z(x) + (z' - z(x))), y) \\ &\approx \sup_{\|z' - z(x)\| \leq \epsilon} \left[\mathcal{L}(f_L^\theta(z(x)), y) + \nabla_z \mathcal{L}(f_L^\theta(z(x)), y)^\top (z' - z(x)) \right] \\ &= \mathcal{L}(f^\theta(x), y) + \sup_{\|z' - z(x)\| \leq \epsilon} \nabla_z \mathcal{L}(f_L^\theta(z(x)), y)^\top (z' - z(x)) \\ &= \mathcal{L}(f^\theta(x), y) + \epsilon \|\nabla_z \mathcal{L}(f_L^\theta(z(x)), y)\|_* , \end{aligned}$$

where $\|\cdot\|_*$ denotes the dual norm corresponding to $\|\cdot\|$. This added regularization term penalizes large gradients in the latent space, promoting robustness by ensuring that small perturbations in $z(x)$ do not lead to significant changes in the loss. By minimizing this term alongside the original loss, the model gains stability and improved performance under real-world distributional shifts.

D DATASET DETAILS

D.1 ORIGINAL DATASET

Colored MNIST (CMNIST) (Arjovsky et al., 2019). The CMNIST dataset is designed for a noisy digit recognition task, incorporating color as a spurious attribute. The dataset is divided into four distinct groups based on class and color: $g_1 = \{0, \text{green}\}$, $g_2 = \{1, \text{green}\}$, $g_3 = \{0, \text{red}\}$, and $g_4 = \{1, \text{red}\}$. It involves two classes: class 0 includes the digits (0, 1, 2, 3, 4) and class 1 includes the digits (5, 6, 7, 8, 9). The training set consists of 30,000 samples, where for class 0, the ratio of red to green samples is 8 : 2, while for class 1, this ratio is 2 : 8. The validation set, which comprises 10,000 samples, maintains an equal distribution of color across both classes, with a 1 : 1 ratio of red to green samples for each class. The test set includes 20,000 samples and introduces a more pronounced group distribution shift: class 0 has a red to green sample ratio of 1 : 9, and class 1 has a ratio of 9 : 1. Following the approach proposed by Arjovsky et al. (2019), labels in the dataset are flipped with a probability of 0.25.

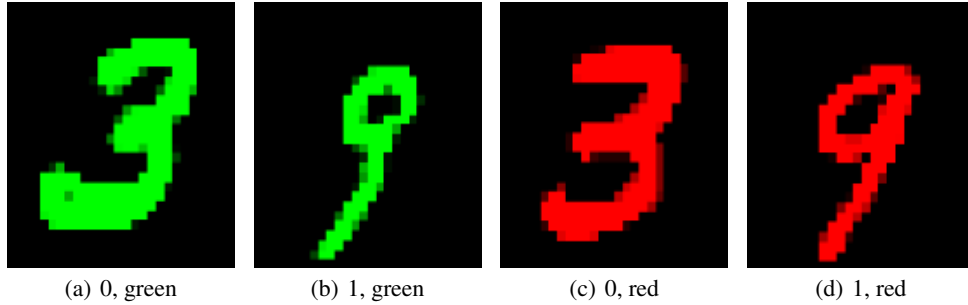


Figure 2: Example images from the CMNIST dataset. The groups are $g_1 = \{0, \text{green}\}$, $g_2 = \{1, \text{green}\}$, $g_3 = \{0, \text{red}\}$, and $g_4 = \{1, \text{red}\}$.

Waterbirds (Sagawa et al., 2020). The Waterbirds dataset is designed to classify images of birds into two categories: “waterbirds” and “landbirds”, with a deliberate introduction of spurious correlations between the bird type and the background. The dataset is divided into four distinct groups based on bird type and background: $g_1 = \{\text{landbird}, \text{land}\}$, $g_2 = \{\text{landbird}, \text{water}\}$, $g_3 = \{\text{waterbird}, \text{land}\}$, and $g_4 = \{\text{waterbird}, \text{water}\}$. This synthetic dataset is created by combining bird images from the

Caltech-UCSD Birds 200-2011 (CUB) dataset (Wah et al., 2011) with backgrounds from the Places dataset (Zhou et al., 2017). Waterbird species, such as albatross, auklet, cormorant, frigatebird, and others, are grouped together, while all other species are classified as landbirds. The dataset comprises 4,795 training samples distributed as follows: 3,498 landbirds on land backgrounds, 1,057 waterbirds on water backgrounds, 184 landbirds on water backgrounds, and 56 waterbirds on land backgrounds. This setup highlights the minority groups and the inherent spurious correlations. In contrast to the training set, the validation and test sets are constructed to have an equal number of samples for each group within each class. The minority group, waterbirds on land, emphasizes the skewed distribution of the dataset, making it suitable for studying the impact of spurious correlations on model performance. The Waterbirds dataset is accessible through the Wilds library in PyTorch (Koh et al., 2021).

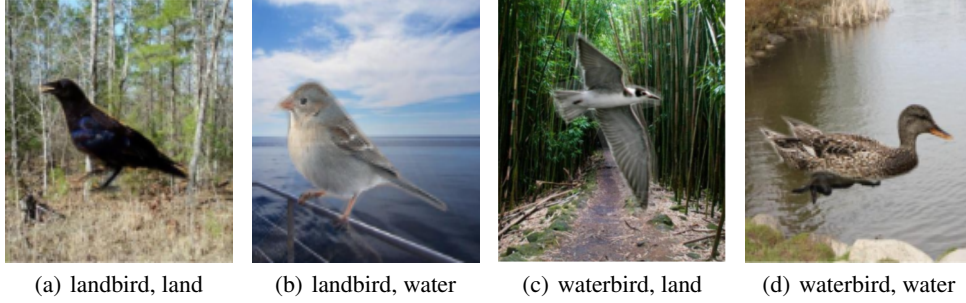


Figure 3: Example images from the Waterbirds dataset. The groups are $g_1 = \{\text{landbird, land}\}$, $g_2 = \{\text{landbird, water}\}$, $g_3 = \{\text{waterbird, land}\}$, and $g_4 = \{\text{waterbird, water}\}$.

CelebA (Liu et al., 2015). The CelebA dataset is used for a hair-color prediction task with facial images of celebrities, where the target labels are “blond” and “non-blond” hair colors. For experimental purposes, the dataset is divided into four distinct groups based on hair color and gender: $g_1 = \{\text{non-blond hair, female}\}$, $g_2 = \{\text{non-blond hair, male}\}$, $g_3 = \{\text{blond hair, female}\}$, and $g_4 = \{\text{blond hair, male}\}$. Gender serves as a spurious feature, introducing correlations between the hair color and gender of individuals. The training set consists of 162,770 images distributed as follows: 71,629 females with non-blond hair, 66,874 males with non-blond hair, 22,880 females with blond hair, and 1,387 males with blond hair. The validation set includes 19,867 images, with 8,535 females with non-blond hair, 8,276 males with non-blond hair, 2,874 females with blond hair, and 182 males with blond hair. The test set comprises 19,962 images, with 9,767 females with non-blond hair, 7,535 males with non-blond hair, 2,480 females with blond hair, and 180 males with blond hair. The minority group in this dataset is males with blond hair, which constitutes a small fraction of the data, highlighting the skewed distribution and the presence of spurious correlations.

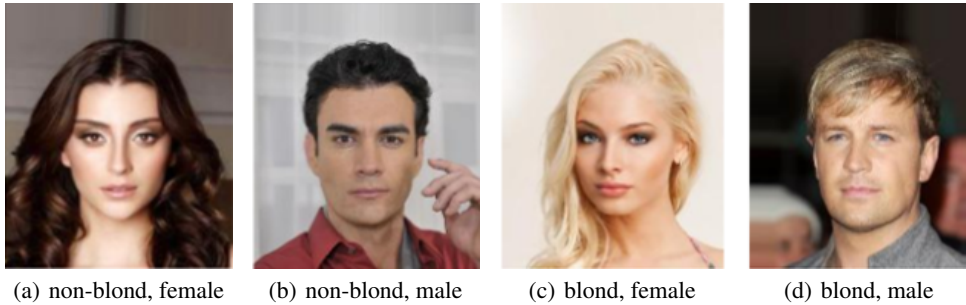


Figure 4: Example images from the CelebA dataset. The groups are $g_1 = \{\text{non-blond hair, female}\}$, $g_2 = \{\text{non-blond hair, male}\}$, $g_3 = \{\text{blond hair, female}\}$, and $g_4 = \{\text{blond hair, male}\}$.

D.2 MODIFIED DATASETS

Building on the previously introduced datasets—CMNIST, Waterbirds, and CelebA—we constructed modified versions of these datasets by applying conditional distribution shifts to the minority groups, simulating real-world scenarios. Below, we detail the modifications for each dataset and illustrate these shifts with corresponding figures.

Modified CMNIST. In the CMNIST dataset, we created a modified version where the minority group’s images (label 1, red) were rotated by 90 degrees in the test set, while they remained unrotated in the training set. This manipulation simulates conditional distribution shifts often encountered in real-world applications. Figure 5 provides an illustration of this shift, showing example images from the train and test sets.

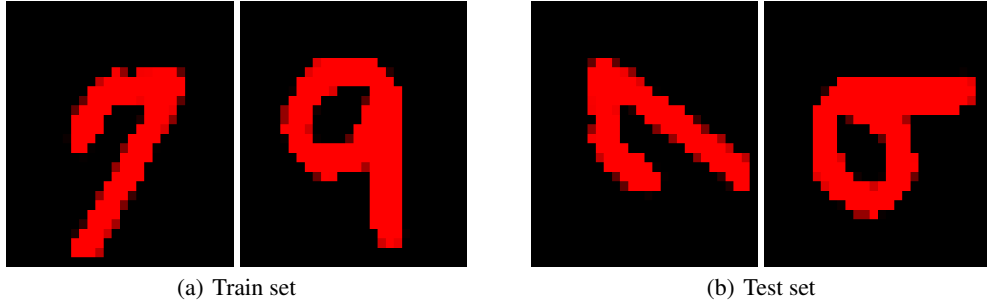


Figure 5: Example of conditional distribution shift in the CMNIST dataset, where the minority group (label 1, red) images are rotated by 90 degrees in the test set, while they are unrotated in the training set.

Modified Waterbirds. For the Waterbirds dataset, we constructed a modified version where the minority group (waterbird, land background) was designed to have a shift in species composition between the train and test sets. Specifically, the training set included only waterfowl species, such as Gadwall, Grebe, Mallard, Merganser, and Pacific Loon, while the test set contained exclusively seabird species, including Albatross, Auklet, Cormorant, Frigatebird, Fulmar, Gull, Jaeger, Kittiwake, Pelican, Puffin, Tern, and Guillemot. During the dataset construction process, we identified and corrected a mislabeling issue involving three species—Western Wood-Pewee, Eastern Towhee, and Western Meadowlark—which had been incorrectly labeled as waterbirds instead of landbirds (Asgari et al., 2022). Figure 6 illustrates this shift, highlighting the separation of species between the train and test sets.



Figure 6: Example of conditional distribution shift in the Waterbirds dataset, where the minority group (waterbird on land background) consists of waterfowl in the training set and seabirds in the test set.

To further highlight the impact of this modification, Figure 7 compares the original distribution and the modified distribution shift scenarios. In the original dataset (Figure 7(a)), bird species in the minority group are relatively evenly distributed across train, validation, and test sets. However, in

the modified version (Figure 7(b)), the training set contains only waterfowl, while the test set is composed entirely of seabirds, creating a distinct distribution shift.

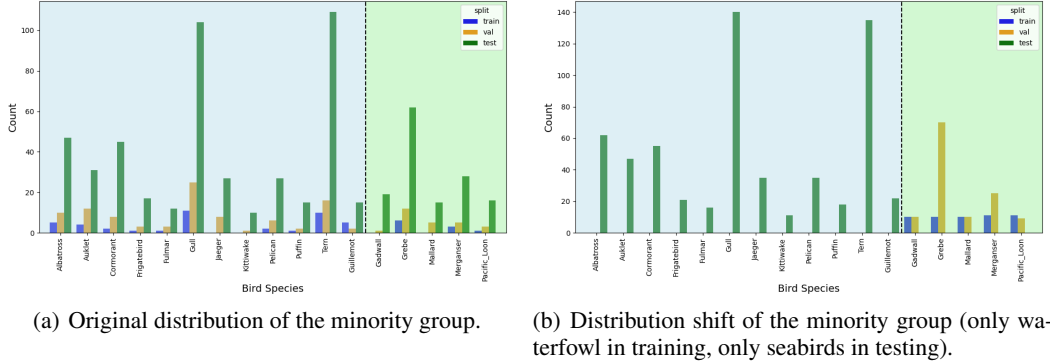


Figure 7: Comparing the original and shifted distributions of the minority group (waterbird, land background) in the Waterbirds dataset (left: seabirds, right: waterfowl, split by dashed line).

Modified CelebA. In the CelebA dataset, we modified the minority group (blond hair, male) to have different attributes between the train and test sets. Specifically, the training set contained only images without glasses, while the test set contained only images with glasses. This modification reflects real-world distribution shifts where rare attributes in small minority groups may change across different distributions, impacting model performance. Figure 8 shows example images demonstrating this shift.

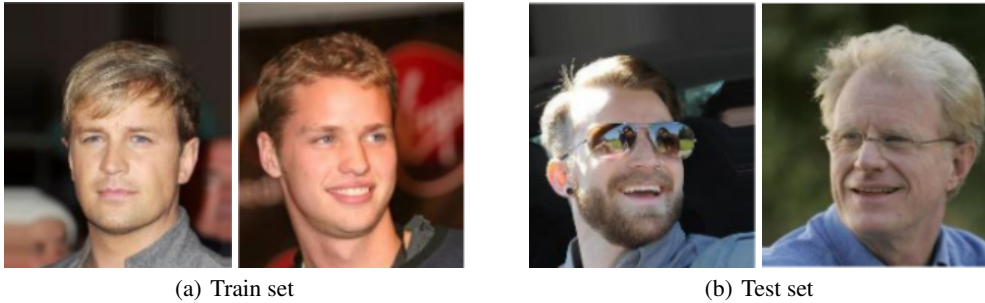


Figure 8: Example of conditional distribution shift in the CelebA dataset, where the minority group (blond hair, male) included only images without glasses in the training set and images with glasses in the test set.

Figure 9 provides a detailed comparison of the original and modified distributions for the CelebA dataset. In the original distribution (Figure 9(a)), the minority group is predominantly represented by the “Without Eyeglasses” category across train, validation, and test sets, with relatively few examples in the “With Eyeglasses” category. In the modified version (Figure 9(b)), the training set consists exclusively of “Without Eyeglasses” images, while the test set contains only “With Eyeglasses”, creating a clear disjoint in key attributes between training and testing phases.

By introducing these conditional distribution shifts, our modified datasets simulate real-world challenges, particularly in scenarios where small minority groups are highly susceptible to such changes. These constructions not only reflect practical settings but also provide realistic benchmarks for evaluating the robustness and generalization capabilities of machine learning models under diverse and challenging conditions.

E BASELINE DETAILS

We compare our method against a range of representative baselines:

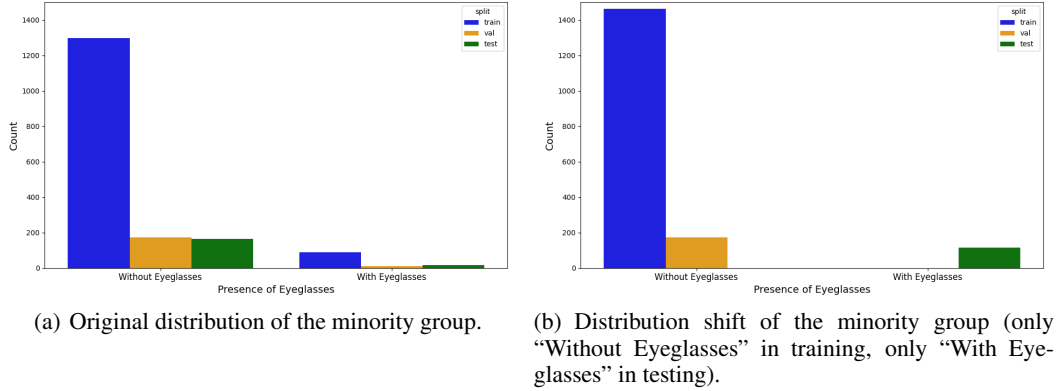


Figure 9: Comparing the original and shifted distributions of the minority group (blond hair, male) in the CelebA dataset (left: “Without Eyeglasses”, right: “With Eyeglasses”).

- **ERM**: ERM optimizes average accuracy on the training set without any robust objective or group-specific considerations.
- **Group DRO** (Sagawa et al., 2020): A canonical approach for mitigating spurious correlations using known group labels. By partitioning data into predefined groups and minimizing the worst-case group loss, Group DRO aims to improve the worst-group accuracy relative to standard ERM.
- **JTT** (Liu et al., 2021): A two-step method that first trains an ERM model to identify misclassified samples (viewed as proxies for minority groups), then upsamples these samples and retrain a classifier.
- **CnC** (Zhang et al., 2022): Identifies samples that share the same true class but differ in spurious attributes by analyzing ERM outputs, then trains a robust model with a contrastive learning objective. This does not require explicit group labels.
- **SSA** (Nam et al., 2022): Infers latent groups via a loss-based criterion, then applies Group DRO to improve robustness. This method partially automates the discovery of group boundaries without needing full group labels.
- **LISA** (Yao et al., 2022): Mitigates spurious correlations by using Mixup strategies. Depending on the dataset, LISA employs different Mixup variants (e.g., classic Mixup, CutMix, Manifold Mix) to interpolate images within the same label or same spurious attribute, thereby reducing reliance on superficial cues.
- **DFR** (Kirichenko et al., 2023): Balances the dataset by subsampling to match the minority group size (the “Subsample” strategy), then retrain an ERM model on this balanced data. This simple yet effective approach can substantially improve worst-group performance. For a fair comparison, following Deng et al. (2024), we evaluate DFR using only the training dataset for both training and fine-tuning, ensuring consistency across methods, which is denoted as DFR^{tr} .
- **PDE** (Deng et al., 2024): Progressively expands the training dataset during the training process, starting with a balanced subset to prevent the model from learning spurious correlations. This approach aims to enhance robustness across all groups, including underrepresented ones.
- **GIC** (Han & Zou, 2024): Uses a two-step pipeline where group membership is partially inferred, then a robust optimization (e.g., Group DRO) is applied. Similar to LISA, it can incorporate tailored Mixup strategies depending on the dataset’s characteristics.

F IMPLEMENTATION DETAILS

For experiments involving our newly constructed datasets, we reimplemented both our proposed method and the relevant baselines. When certain baselines lacked reported results for a given dataset, we used the performance from Zhang et al. (2022) and Han & Zou (2024) if available; otherwise, we performed our own reimplementations under consistent settings. In particular, for the original CMNIST dataset, we reimplemented experiments for DFR and PDE, since their original papers did not

include CMNIST results. In all other cases, we referenced performance metrics from each baseline’s primary source. All experiments were conducted on an NVIDIA GeForce RTX 3090 GPU, and we make our code available at <https://github.com/Sung-Ho-Jo/hierarchical-dro>.

Across all datasets, we employed the torchvision implementation of ResNet-50 pretrained on ImageNet, training with SGD at a momentum of 0.9 and a batch size of 128, following Sagawa et al. (2020). Our approach also introduces a perturbation parameter ϵ to control within-group uncertainty. Specifically, we define $\epsilon_g = \epsilon / \sqrt{n_g}$, where n_g represents the size of group g in the training data. To determine ϵ , we performed a grid search over the set $\{12/255, 24/255, 36/255, 48/255, 60/255, 72/255, 84/255, 96/255\}$, scaling each value by $\sqrt{n_{\min}}$. Here, $n_{\min} = \min_g n_g$ denotes the smallest group size in the training set. Additionally, we tuned the generalization adjustment parameter C over $\{0, 1, 2, 3\}$, as described in Section 3.3 of Sagawa et al. (2020). This setup was applied consistently across every dataset.

For CMNIST, we conducted a grid search over learning rates $\{10^{-4}, 10^{-3}, 10^{-2}\}$ and ℓ_2 penalties $\{10^{-1}, 10^{-2}, 10^{-4}\}$ for 50 epochs. Due to instability in training with the selected parameter combinations in the original Group DRO implementation, we applied a ReduceLROnPlateau scheduler starting at a learning rate of 0.01, using it consistently for both our method and Group DRO to ensure fairness. For Waterbirds, the learning rate was tuned over $\{10^{-3}, 10^{-4}, 10^{-5}\}$ and the ℓ_2 penalty over $\{10^{-4}, 10^{-1}, 1\}$, with training conducted for 300 epochs. For CelebA, the learning rate was tuned over $\{10^{-4}, 10^{-5}\}$ and the ℓ_2 penalty over $\{10^{-4}, 10^{-2}, 1\}$ for 30 epochs. We referred to prior works including Yao et al. (2022) and Ghosal & Li (2023) to guide these hyperparameter search ranges.

G TUNING AND ANALYSIS OF ϵ

G.1 SELECTION OF ϵ

As is common in DRO problems, selecting the size of the ambiguity set ϵ is challenging. Since the extent of distributional shifts is generally unknown, many robust learning methods commonly rely on heuristic or even arbitrary choices in practice. We propose a simple data-driven procedure that partitions the training data using a one-dimensional t-SNE (Van der Maaten & Hinton, 2008) ordering. This avoids reliance on shifted or test-time data while still creating validation splits that mimic intra-group distribution shifts. Our design is inspired by prior work (Duchi & Namkoong, 2021), which used training-data partitioning to simulate distributional shifts.

Specifically, we project each $z(x_i)$ onto a one-dimensional space using t-SNE, rank samples within each group, and split them into five quantiles. The two extreme quantiles (top 20% and bottom 20%) are alternately held out as validation sets, with the remaining 80% used for training. This procedure creates realistic validation shifts that disproportionately affect minority groups.

Finally, we select the value of ϵ that maximizes minority-group accuracy across these validation setups, ensuring that the chosen perturbation radius provides meaningful robustness for underrepresented subpopulations.

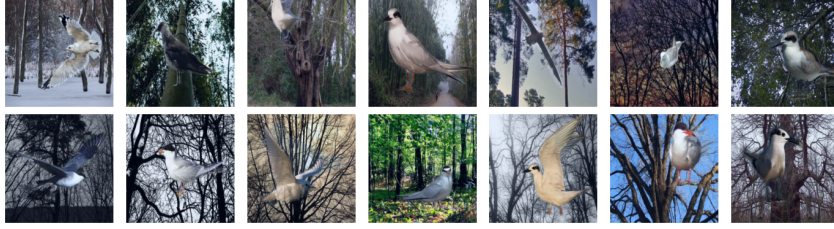
G.2 VISUALIZING T-SNE ORDERING FOR MINORITY GROUPS

To highlight the utility of t-SNE ordering in simulating realistic distribution shifts, we present visual examples of waterbirds from the top 20% and bottom 20% quantiles of the t-SNE projections. This approach effectively partitions samples based on semantically meaningful intra-group differences, enabling validation splits that closely mimic real-world distribution shifts.

As shown in Figure 10, the t-SNE ordering captures nuanced intra-group differences within this minority group. In the top 20% quantile (Figure 10(a)), waterbirds with longer beaks dominate; in the bottom 20% quantile (Figure 10(b)), shorter-beaked waterbirds are more prevalent. This contrast illustrates how a t-SNE-driven partition can create validation splits that mimic real-world distribution shifts. This method not only emphasizes variations within groups but also systematically evaluates the model’s robustness under challenging real-world conditions.



(a) Example images from the top 20% of t-SNE ordering

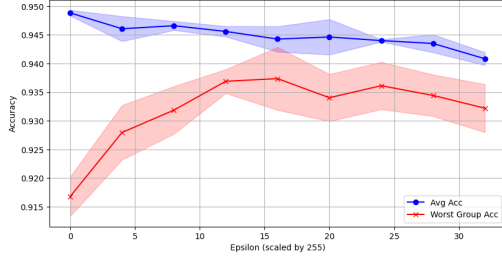


(b) Example images from the bottom 20% of t-SNE ordering

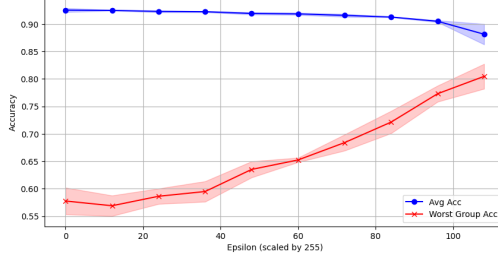
Figure 10: t-SNE-based ordering reveals subtle distinctions within the minority group. (a) The top quantile features waterbirds with longer beaks, while (b) the bottom quantile features those with shorter beaks.

G.3 IMPACT OF ϵ ON ROBUSTNESS

The perturbation parameter ϵ plays a critical role in improving robustness under minority group shifts. Figures 11(a) and 11(b) show how increasing ϵ affects worst-group accuracy for the Waterbirds and CelebA datasets, respectively. Notably, both datasets achieve significant gains in worst-group accuracy when ϵ is set above zero, indicating enhanced resilience to distributional shifts.



(a) Waterbirds dataset under a minority group shift.



(b) CelebA dataset under a minority group shift.

Figure 11: Impact of ϵ on robustness. The x-axis represents ϵ values scaled by 255, and the y-axis indicates accuracy. Each point is the mean of 3 runs (solid lines), and the shaded regions show the standard deviation. For this analysis, the learning rate and ℓ_2 penalty were fixed to isolate the effect of ϵ .

As illustrated in Figure 11(a), larger ϵ values consistently improve worst-group accuracy on Waterbirds, enabling the model to better manage intra-group variations and subpopulation shifts. A similar trend appears in Figure 11(b) for CelebA, further validating the robustness gained by appropriately increasing ϵ .

These findings underscore the importance of incorporating conditional distribution uncertainty into the training framework. By effectively capturing within-group variability, our approach significantly enhances worst-group performance, making it well-suited for handling realistic distributional shifts.

H GRAD-CAM RESULTS AND ANALYSIS

To gain further insight into where each model focuses its attention under minority-group shifts, we visualize Grad-CAM (Selvaraju et al., 2017) heatmaps on misclassified examples (by Group DRO) that our method classifies correctly. Figure 12 shows examples on the Waterbirds dataset, while Figure 13 presents examples from CelebA.

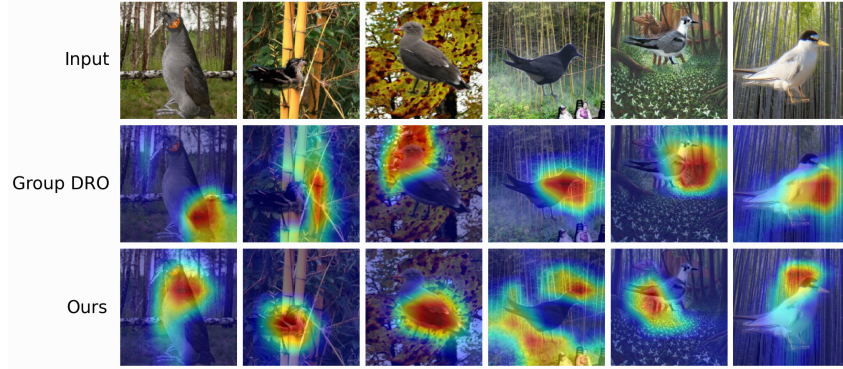


Figure 12: Grad-CAM visualizations for Waterbirds test images from a minority-group shift scenario. Each column shows an input image (top row), Grad-CAM for Group DRO (middle row), and Grad-CAM for our method (bottom row).

Waterbirds. In Figure 12, the minority-group shift involves species changes not observed in the training set. While Group DRO often localizes on a narrow region of the bird—sometimes near the torso or background—our method exhibits a more distributed attention, covering details like the wings, beak, or feet. This broader localization helps the model rely on features invariant to previously unseen waterbird species, enabling robust classification despite changes in the specific types of waterbirds encountered.

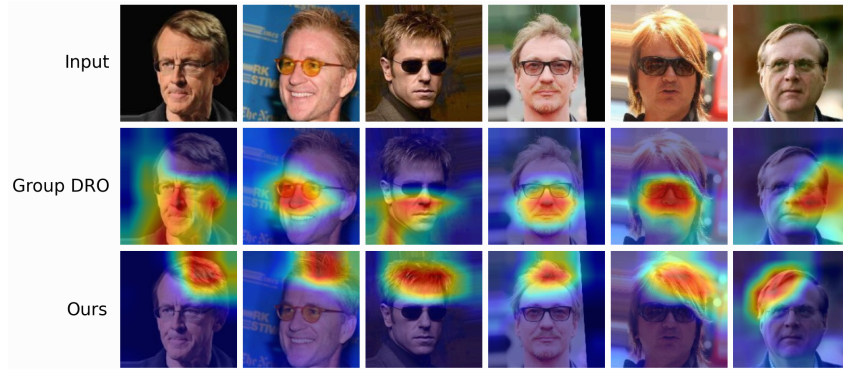


Figure 13: Grad-CAM visualizations for CelebA test images from the minority-group shift scenario. Each column shows the input image (top row), Grad-CAM for Group DRO (middle row), and Grad-CAM for our method (bottom row).

CelebA. Figure 13 shows examples from the minority group (blond-hair, male) in which the test images include glasses—an attribute absent from the training set. In these cases, Group DRO erroneously directs attention toward the facial or eyewear regions rather than focusing on hair color. By contrast, our method more reliably highlights the hair region, aligning with the intended classification objective and enabling correct predictions even under previously unseen attributes.

Overall, these visualizations confirm that, under challenging distribution shifts, our hierarchical DRO framework is less prone to confounding features and more successful in focusing on the task-

relevant regions. This broader and more contextually aligned attention helps maintain strong performance even when encountering unseen or spurious attributes.

I COMPARISON WITH STANDARD WASSERSTEIN DRO

To contextualize the performance gains achieved by our hierarchical formulation, we compare it against standard Wasserstein DRO, using the same latent-space cost function and identical training hyperparameters. The comparison is carried out on the Waterbirds and CelebA benchmarks under the minority group shift setting considered in our experiments, which reflects the primary distributional scenario targeted by our method. Figures 14(a) and 14(b) summarize the behavior of standard Wasserstein DRO across a wide range of Wasserstein radii.

Two structural limitations underlie the weak performance observed in these plots. First, standard Wasserstein DRO models uncertainty only through perturbations within a single Wasserstein ball around the empirical distribution, without incorporating any uncertainty over group proportions. In spurious-correlation settings, however, variation in group proportions constitutes a fundamental mode of distribution shift: a minority group that is scarcely observed during training may appear far more frequently at test time, and ignoring this possibility prevents standard Wasserstein DRO from addressing such mixture-level changes. Second, in spurious-correlation benchmarks, different groups have substantially different sample sizes. Our method is designed to reflect this by assigning a distinct within-group radius to each group ($\epsilon_g = \epsilon/\sqrt{n_g}$), so that smaller groups—whose empirical distributions are statistically less reliable—are assigned a larger perturbation radius. In contrast, standard Wasserstein DRO uses a single global radius and therefore cannot adapt to group-specific statistical uncertainty, which is crucial in the minority-shift setting.

For both Waterbirds and CelebA, we varied the Wasserstein radius over a broad range of values and report the corresponding worst-group accuracy curves below. Across all radii, standard Wasserstein DRO achieves significantly lower worst-group accuracy than our method, which attains 93.7% and 72.1% on Waterbirds and CelebA, respectively. These results highlight that perturbing the unconditional distribution alone is insufficient for capturing the types of shifts present in spurious-correlation settings.

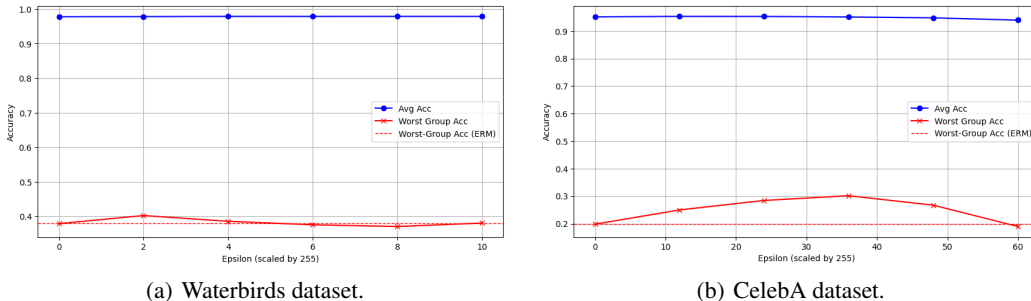


Figure 14: Performance of standard Wasserstein DRO across a wide range of radii.

These findings reinforce the main conclusion of our study: achieving robustness under the complex and realistic distribution shifts that arise in spurious-correlation benchmarks—particularly those involving minority-group shifts—requires modeling uncertainty over group proportions as well as group-specific conditional distributions.