

Generative Neural Networks for Sinkhorn Distributionally Robust Hypothesis Testing

Fenglin Zhang, Teyan Liu, Jie Wang*

August 20, 2026

Abstract

This paper studies the Sinkhorn distributionally robust hypothesis testing (SDRHT) problem, seeking a robust detector against least-favorable distributions in Sinkhorn discrepancy-based ambiguity sets centered at the empirical distributions. Existing approaches solve this problem by solving large-scale conic programs, which are not scalable. To overcome this, we propose a generative framework that learns least-favorable distributions and supports efficient training and end-to-end sampling. For the Sinkhorn discrepancy-based ambiguity sets, we first derive an equivalent conditional-KL-divergence representation with respect to kernel-smoothed reference distributions. This property allows us to prove strong duality for both constrained and unconstrained minimax SDRHT formulations. Based on the closed-form optimal detector and Brenier’s theorem, we reformulate the max-min dual formulation as a maximization problem over convex potentials whose gradients characterize invertible transport maps between kernel-smoothed distributions and their least-favorable counterparts. We efficiently approximate these potentials using Hyper Input Convex Neural Networks (HyCNNs) equipped with stochastic gradient estimators and prove the representation power of HyCNNs and the distributional universality of their induced transport maps. Numerical results show that the proposed method achieves superior accuracy and robustness across different sample sizes and dimensions, while avoiding the scalability limitations of classical SDRHT methods.

Keywords: Sinkhorn distributionally robust optimization, robust hypothesis testing, generative model, least-favorable distribution generation, hyper input convex neural network

1 Introduction

Hypothesis testing is a fundamental problem in statistics and scientific discovery, aiming to find an optimal detector that, given new data, can discriminate between two competing hypotheses while minimizing decision errors. The Neyman–Pearson lemma [1] yields the optimal likelihood-ratio test when the densities of the underlying distributions are known. In practice, these distributions are typically unknown; instead, the decision-maker has only access to their independent

*Fenglin Zhang is affiliated with School of Artificial Intelligence, The Chinese University of Hong Kong, Shenzhen; Teyan Liu is affiliated with Department of Decision Analytics and Operations, College of Business, City University of Hong Kong; Jie Wang is affiliated with School of Artificial Intelligence and School of Data Science, The Chinese University of Hong Kong, Shenzhen. Corresponding author: Jie Wang, jwang@cuhk.edu.cn.

and identically distributed (i.i.d.) samples. This setting motivates data-driven hypothesis tests that are both statistically reliable and computationally efficient.

However, data-driven tests trained on finite samples are vulnerable to distributional misspecification [2]. On the one hand, for some applications like healthcare [3] and anomaly detection [4], where samples are limited and noisy, reliable estimation of the underlying distribution is challenging due to inevitable measurement and statistical errors. On the other hand, real data may be non-stationary or may undergo distributional shifts, leading to differences between training and testing data. To improve the reliability of hypothesis testing, especially in high-stakes scenarios, robust methods are necessary to circumvent these challenges [5, 6].

Distributionally robust optimization (DRO) [7] offers a powerful paradigm for making robust decisions under distributional uncertainty. When applied to hypothesis testing, DRO optimizes the robust detector against the least-favorable distributions (also known as worst-case distributions) within an ambiguity set consisting of plausible distributions, safeguarding against misspecification and enhancing out-of-sample reliability when applied to new, unseen data [8]. The design of the ambiguity set is crucial for balancing computational tractability and robust out-of-sample performance [9]. Popular choices include Huber’s contamination model [10], f -divergence [11, 12], moments [13, 14], kernels [15, 16], or Wasserstein distance [8, 17].

DRO with ambiguity sets constructed using the Sinkhorn discrepancy has recently received increasing attention [18, 19]. Conventional Wasserstein DRO-based hypothesis testing restricts the worst-case distributions to have the same support as that of the empirical distributions [8]. By adding entropic regularization to the Wasserstein distance, the Sinkhorn discrepancy provides a smooth alternative to the Wasserstein distance [20], and the Sinkhorn discrepancy-based ambiguity sets naturally encourage continuous distributional estimates. Thus, unlike the Wasserstein distance-based ambiguity set, if a testing sample is outside the support of the observed data, it is still feasible to leverage the likelihood ratio detector to design a robust test. However, finding the least-favorable distributions in Sinkhorn-based hypothesis testing is an infinite-dimensional optimization problem. Existing approaches approximately solve it using sample average approximation and result in large-scale conic programs [21, 22], whose computational cost grows rapidly with the sample size. It is desirable to develop more tractable approaches for solving the robust hypothesis testing problem with Sinkhorn-based ambiguity sets.

Parallel to these developments in robust statistics, flow-based generative models have shown remarkable success in learning complex distributions as pushforwards of simple reference distributions through invertible and differentiable maps [23, 24, 25]. This framework is not only adept at generating new samples but is also naturally suited for solving optimization problems over probability space, e.g., generating least-favorable distributions in DRO [26, 27, 28]. By parameterizing a distribution as a flow-based transformation starting from a probability density, one can efficiently optimize an objective over the probability space using gradient-based methods [29, 30].

In this work, we leverage flow-based generative models to solve the Sinkhorn distributionally robust hypothesis testing (SDRHT). Based on the theoretical results of optimal transport, we introduce a generative framework to learn continuous least-favorable distributions within the Sinkhorn-based ambiguity sets by optimizing convex potentials. By parameterizing the convex

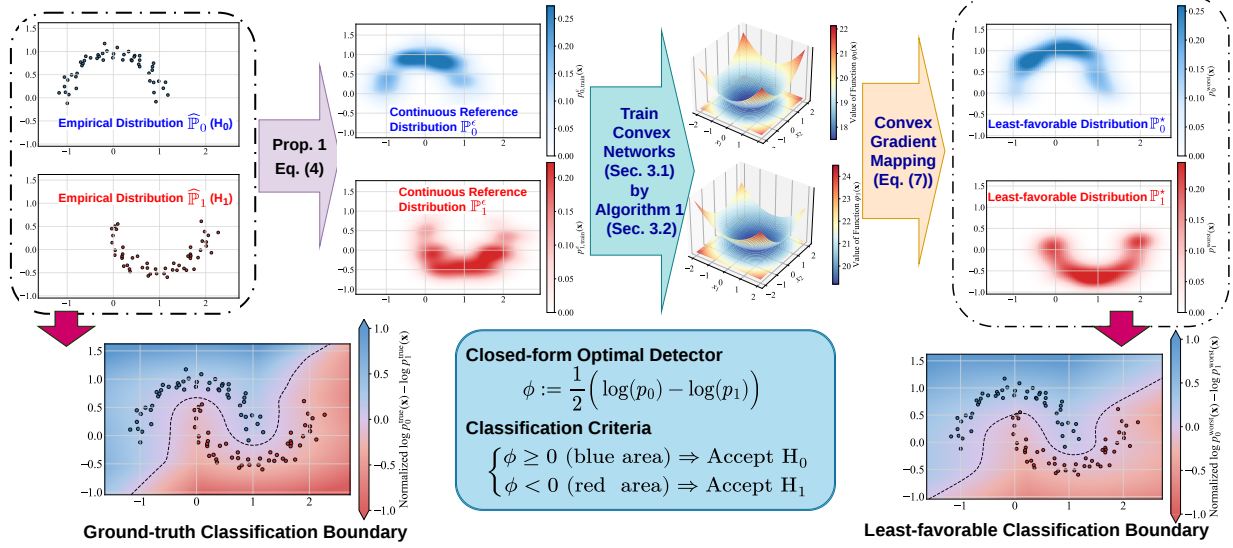


Figure 1: Illustration of the training and testing processes of our least-favorable detector.

potentials, our approach converts the original infinite-dimensional optimization into a finite-dimensional and tractable problem. Our main contributions are summarized as follows:

- **Novel Reformulation:** We establish a strong duality result for SDRHT and thereby reformulate the original minimax optimization as a max-min problem. As the inner minimized detector can be explicitly derived, this problem is equivalent to a maximum problem over distributions. Leveraging Brenier’s theorem, we further reformulate it as a maximization problem over convex potentials.
- **Generative Framework:** We parameterize the convex potentials by Hyper Input Convex Neural Networks (HyCNNs) and develop stochastic gradient estimators for the objective. This transforms the infinite-dimensional optimization into a finite-dimensional one, enabling efficient and scalable training and sampling. We also prove the representation power of HyCNNs and the distributional universality of their induced gradient maps.
- **Empirical Validation:** We validate the performance of our method through extensive numerical studies on both synthetic datasets and real-world data with different sample sizes and dimensions, demonstrating its superiority over existing baselines in terms of both accuracy and robustness.

Figure 1 illustrates the training and testing processes of the proposed method using a toy Moons dataset. By exploiting the geometry of the Sinkhorn-based ambiguity set (Proposition 1 and Eq. (4)), the least-favorable distributions ($\mathbb{P}_0^*, \mathbb{P}_1^*$) are generated by first applying kernel smoothing to the empirical distributions $\widehat{\mathbb{P}}_0$ and $\widehat{\mathbb{P}}_1$ to obtain $\mathbb{P}_0^\epsilon$ and $\mathbb{P}_1^\epsilon$, and next robustifying them. In the robustification procedure, we learn convex-gradient maps (see Eq. (7) and Algorithm 1) that generate $\mathbb{P}_0^*$ and $\mathbb{P}_1^*$ based on the kernel-smoothed reference distributions $\mathbb{P}_0^\epsilon, \mathbb{P}_1^\epsilon$. The resulting least-favorable distributions induce a nonlinear robust decision boundary with high accuracy.

The remainder of this paper is organized as follows. Section 2 presents the models and properties of SDRHT. Section 3 develops a framework for generating least-favorable distributions. Section 4 shows the representation power of HyCNNs. Section 5 reports the numerical results. Section 6 concludes this paper. Proofs and additional technical notes are provided in the online appendix to this paper.

Notations. For $K \in \mathbb{N}$, let $[K] := \{1, \dots, K\}$. We write $\mathbb{P} \ll \nu$ if the probability measure $\mathbb{P}$ is absolutely continuous with respect to the measure ν . We denote $\mathbb{P} \otimes \mathbb{Q}$ by the product measure of two probability measures $\mathbb{P}$ and $\mathbb{Q}$. For a transport map $T : \Omega \rightarrow \Omega$, the pushforward of a probability measure $\mathbb{P}$ is denoted as $T_{\#}\mathbb{P}$, such that $T_{\#}\mathbb{P}(A) = \mathbb{P}(T^{-1}(A))$ for any measurable set $A \subseteq \Omega$. For a convex function $\varphi : \mathbb{R}^d \rightarrow \mathbb{R} \cup \{\infty\}$, denote its convex conjugate by $\varphi^*(\mathbf{y}) = \sup_{\mathbf{x} \in \mathbb{R}^d} \{\mathbf{y}^\top \mathbf{x} - \varphi(\mathbf{x})\}$ for any $\mathbf{y} \in \mathbb{R}^d$.

2 Sinkhorn Distributionally Robust Hypothesis Testing

In this section, we introduce the Sinkhorn distributionally robust hypothesis testing (SDRHT), derive its strong duality property, and obtain its equivalent functional optimization reformulation.

Let $\Omega \subseteq \mathbb{R}^d$ be a closed sample space. Denote $\mathcal{P}(\Omega)$ as the set of all probability measures on Ω , and $\mathcal{M}(\Omega)$ as the set of measures on Ω . Denote by $\mathcal{P}_k \subset \mathcal{P}(\Omega)$ the ambiguity set under hypotheses H_k , where $k \in \mathbb{K} = \{0, 1\}$. Given a testing sample ω , the composite hypothesis testing decides which one of the following hypotheses is true:

$$H_0 : \omega \sim \mathbb{P}_0, \mathbb{P}_0 \in \mathcal{P}_0; \quad H_1 : \omega \sim \mathbb{P}_1, \mathbb{P}_1 \in \mathcal{P}_1.$$

A hypothesis test (detector) $\phi : \Omega \rightarrow \mathbb{R}$ is a measurable function. For any testing sample ω , it accepts the null hypothesis H_0 and rejects alternative hypothesis H_1 whenever $\phi(\omega) \geq 0$, otherwise accepts H_1 and rejects H_0 . The exact risk of this detector is defined as the sum of type-I and type-II errors. For this non-convex objective, we follow the generating-function approach [8, 31] and optimize the following convex surrogate risk instead:

$$\Phi(\phi; \mathbb{P}_0, \mathbb{P}_1) := \mathbb{E}_{\omega \sim \mathbb{P}_0}[\ell \circ (-\phi)(\omega)] + \mathbb{E}_{\omega \sim \mathbb{P}_1}[\ell \circ \phi(\omega)],$$

where the generating function ℓ is defined as follows.

Definition 1 (Generating Function). *A generating function $\ell : \mathbb{R} \rightarrow \mathbb{R}_+ \cup \{\infty\}$ is a non-negative, non-decreasing, convex function such that $\ell(0) = 1$ and $\lim_{t \rightarrow -\infty} \ell(t) = 0$.* $\diamond$

For some common choices of the generating function, Table 1 summarizes the corresponding closed-form optimal detectors ϕ_{opt} given distributions $\mathbb{P}_0$ and $\mathbb{P}_1$ [8, 21], where p_0 and p_1 represent the probability density functions of $\mathbb{P}_0$ and $\mathbb{P}_1$. In this paper, we choose the generating function $\ell(t) := \exp(t)$.

Given sets of training samples $\{\mathbf{x}_k^{(1)}, \dots, \mathbf{x}_k^{(n_k)}\}$ for all $k \in \mathbb{K}$, denote the corresponding empirical distributions as $\widehat{\mathbb{P}}_k = \frac{1}{n_k} \sum_{i=1}^{n_k} \delta_{\mathbf{x}_k^{(i)}}$, where n_k denotes the number of samples from $\mathbb{P}_k \in \mathcal{P}_k$. The following definition introduces the Sinkhorn discrepancy, an entropy-regularized variant of the Wasserstein distance that quantifies the distance between distributions [18, 21]. Throughout the paper, we use the quadratic transport cost $c(\mathbf{x}, \mathbf{z}) := \frac{1}{2} \|\mathbf{x} - \mathbf{z}\|_2^2$.

Table 1: Common choices of generating function (first column), together with their corresponding optimal detector (second column) and surrogate risk (third column).

$\ell(t)$	Optimal ϕ	Corresponding optimal $1 - \Phi(\mathbb{P}_0, \mathbb{P}_1)/2$
$\exp(t)$	$\log \sqrt{p_0/p_1}$	$H^2(\mathbb{P}_0, \mathbb{P}_1)$
$\log(1 + \exp(t))/\log 2$	$\log(p_0/p_1)$	$\text{JS}(\mathbb{P}_0, \mathbb{P}_1)/\log 2$
$(t+1)^2$	$1 - 2p_0/(p_0 + p_1)$	$\chi^2(\mathbb{P}_0, \mathbb{P}_1)$
$(t+1)_+$	$\text{sign}(p_0 - p_1)$	$\text{TV}(\mathbb{P}_0, \mathbb{P}_1)$

H^2 : Hellinger divergence; JS: Jensen-Shannon divergence; TV: Total variation distance.

Definition 2 (Sinkhorn Discrepancy). *Consider any two distributions $\mathbb{P}, \mathbb{Q} \in \mathcal{P}(\Omega)$ and reference measure $\nu \in \mathcal{M}(\Omega)$ such that $\mathbb{Q} \ll \nu$. For regularization parameter $\epsilon > 0$, the Sinkhorn discrepancy between two distributions $\mathbb{P}$ and $\mathbb{Q}$ is defined as*

$$\mathcal{S}_\epsilon(\mathbb{P}, \mathbb{Q}) = \inf_{\gamma \in \Gamma(\mathbb{P}, \mathbb{Q})} \left\{ \mathbb{E}_{(\mathbf{x}, \mathbf{z}) \sim \gamma} [c(\mathbf{x}, \mathbf{z})] + \epsilon \cdot \mathcal{D}_{\text{KL}}(\gamma \| \mathbb{P} \otimes \nu) \right\},$$

where $c(\mathbf{x}, \mathbf{z})$ represents the transport cost function, $\Gamma(\mathbb{P}, \mathbb{Q})$ denotes the set of couplings of $\mathbb{P}$ and $\mathbb{Q}$, respectively, and $\mathcal{D}_{\text{KL}}(\gamma \| \mathbb{P} \otimes \nu)$ denotes the relative entropy of distribution γ with respect to measure $\mathbb{P} \otimes \nu$, i.e.,

$$\mathcal{D}_{\text{KL}}(\gamma \| \mathbb{P} \otimes \nu) = \int_{\Omega \times \Omega} \log \frac{d\gamma(\mathbf{x}, \mathbf{z})}{d\mathbb{P}(\mathbf{x})d\nu(\mathbf{z})} d\gamma(\mathbf{x}, \mathbf{z}). \quad \diamond$$

Let the space of detectors be $\mathcal{F}$. We consider the following SDRHT problem, optimizing a detector to minimize the least-favorable risk in ambiguity sets:

$$\inf_{\phi \in \mathcal{F}} \sup_{\mathbb{P}_0 \in \mathbb{B}_{\rho_0}(\widehat{\mathbb{P}}_0), \mathbb{P}_1 \in \mathbb{B}_{\rho_1}(\widehat{\mathbb{P}}_1)} \Phi(\phi; \mathbb{P}_0, \mathbb{P}_1), \quad (\text{SDRHT})$$

where the ambiguity sets are constructed by the Sinkhorn discrepancy $\mathcal{S}_\epsilon(\widehat{\mathbb{P}}_k, \mathbb{P})$, offering a flexible, non-parametric framework to characterize the least-favorable distributions

$$\mathbb{B}_{\rho_k}(\widehat{\mathbb{P}}_k) := \{\mathbb{P} : \mathcal{S}_\epsilon(\widehat{\mathbb{P}}_k, \mathbb{P}) \leq \rho_k\}, \quad \forall k \in \mathbb{K}. \quad (1)$$

We also consider the following soft-constrained counterpart of problem (SDRHT):

$$\inf_{\phi \in \mathcal{F}} \sup_{\mathbb{P}_0, \mathbb{P}_1 \in \mathcal{P}(\Omega)} \Phi(\phi; \mathbb{P}_0, \mathbb{P}_1) - \sum_{k \in \mathbb{K}} \lambda_k \mathcal{S}_\epsilon(\widehat{\mathbb{P}}_k, \mathbb{P}_k), \quad (\text{soft-SDRHT})$$

where $\lambda_k > 0$ for all $k \in \mathbb{K}$ are penalty parameters.

2.1 Strong Duality

In this subsection, since the optimal detector admits a closed form for fixed distributions, we prove the strong duality of problems (SDRHT) and (soft-SDRHT) by interchanging the minimization over detectors and the maximization over distributions. Note that the functional space $\mathcal{F}$ is infinite-dimensional and naturally convex, but is generally not compact. Hence, the classical

Sion's minimax theorem cannot be used to verify strong duality. Instead, we prove strong duality primarily using the lopsided minimax theorem [7], which does not require $\mathcal{F}$ to be compact.

For the transport cost and Sinkhorn discrepancy, we introduce the following assumptions.

- Assumption 1.** (I) *The transport cost function $c(x, z)$ is $\widehat{\mathbb{P}} \otimes \nu$ -measurable and satisfies $\nu\{z : 0 \leq c(x, z) < \infty\} = 1$ for $\widehat{\mathbb{P}}$ -almost every x .*
- (II) $\mathbb{E}_{z \sim \nu}[e^{-c(x, z)/\epsilon}] < \infty$ for $\widehat{\mathbb{P}}$ -almost every x .
- (III) *The cost function $c(x, z)$ is lower semicontinuous in (x, z) .*
- (IV) *For every joint distribution γ on $\Omega \times \Omega$ with first marginal distribution $\mathbb{P}$, it has a regular conditional distribution γ_x given the value of the first marginal equals x .*

According to [18], Assumptions 1(I) and (II) ensure that both the Sinkhorn discrepancy and the equivalent reformulation in the following Proposition 1 are well-defined. Assumption 1(III) ensures that the Sinkhorn discrepancy is weakly lower semicontinuous. Assumption 1(IV) ensures that the joint distribution γ can be written as $d\gamma(x, z) = d\widehat{\mathbb{P}}(x)d\gamma_x$ and the law of total expectation holds. Our choice of transport cost $c(x, z) = \frac{1}{2}\|x - z\|_2^2$ naturally satisfies Assumption 1(III).

Define the non-negative loss function

$$L_k(\phi, \omega) := \ell\left((-1)^{k+1} \cdot \phi(\omega)\right), \quad \forall k \in \mathbb{K}. \quad (2)$$

In this paper, we focus on continuous detectors and consider the following assumptions for the loss function.

- Assumption 2** (Loss Function). (I) *For any $k \in \mathbb{K}$, $\phi \in \mathcal{F}$, and $\beta_k > 0$, the reference distribution $\mathbb{P}_k^\epsilon$ in Lemma 1 satisfies $\mathbb{E}_{\omega \sim \mathbb{P}_k^\epsilon}[\exp(\beta_k \cdot L_k(\phi, \omega))] < \infty$.*
- (II) *There exists $\mathbb{Q} \in \mathcal{P}(\Omega)$ with $\inf_{\phi \in \mathcal{F}} \mathbb{E}_{\omega \sim \mathbb{Q}}[L_k(\phi, \omega) - \lambda_k \mathcal{S}_\epsilon(\widehat{\mathbb{P}}_k, \mathbb{Q})] > -\infty$ for any $k \in \mathbb{K}$.*

Assumption 2(I) leads to the following Lemma 2. Assumption 2(II) ensures that the optimal value of the dual of problem (soft-SDRHT) is strictly larger than $-\infty$.

The following proposition provides an equivalent conditional-KL-divergence representation of the Sinkhorn-based ambiguity set.

Proposition 1 (Conditional-KL Representation of the Sinkhorn Ambiguity Set). *The ambiguity sets based on Sinkhorn discrepancy with $\epsilon > 0$ are equivalent to*

$$\mathbb{B}_{\bar{\rho}_k}(\mathbb{P}_k^\epsilon) := \left\{ \mathbb{Q} : \inf_{\gamma \in \Gamma(\widehat{\mathbb{P}}_k, \mathbb{Q})} \mathbb{E}_{x \sim \widehat{\mathbb{P}}_k} \left[\mathcal{D}_{\text{KL}}(\gamma_x \| \mathcal{K}_{x, \epsilon}) \right] \leq \bar{\rho}_k / \epsilon \right\}, \quad \forall k \in \mathbb{K}, \quad (3)$$

where $\bar{\rho}_k := \rho_k + \epsilon \mathbb{E}_{x \sim \widehat{\mathbb{P}}_k} \left[\log \int_{\mathbb{R}^d} e^{-c(x, z)/\epsilon} dz \right]$, probability density $d\mathbb{P}_k^\epsilon(z) := \mathbb{E}_{x \sim \widehat{\mathbb{P}}_k} [d\mathcal{K}_{x, \epsilon}(z)]$,

$$d\mathcal{K}_{x, \epsilon}(z) := \frac{e^{-c(x, z)/\epsilon}}{\int_{\mathbb{R}^d} e^{-c(x, u)/\epsilon} \cdot d\nu(u)} \cdot d\nu(z).$$

The proof of Proposition 1 is provided in Appendix B.1. Proposition 1 shows the connection between the ambiguity set based on Sinkhorn discrepancy and those based on the KL divergence [12]. Since the reference measure $\nu \in \mathcal{M}(\Omega)$, $\mathcal{K}_{x,\epsilon}$ is a probability kernel, and its corresponding continuous reference distributions in Eq. (3) are probability measures, i.e.,

$$\mathbb{P}_k^\epsilon = \frac{1}{n_k} \sum_{i=1}^{n_k} \mathcal{K}_{\mathbf{x}_k^{(i)}, \epsilon} \in \mathcal{P}(\Omega), \quad \forall k \in \mathbb{K}. \quad (4)$$

When $\Omega = \mathbb{R}^d$, ν is a Lebesgue measure, and the transport cost is defined as $\frac{1}{2}\|x - z\|_2^2$, distribution $\mathbb{P}_k^\epsilon$ is a Gaussian mixture because $\mathcal{K}_{\mathbf{x}_k^{(i)}, \epsilon} = \mathcal{N}(\mathbf{x}_k^{(i)}, \epsilon \cdot \mathbf{I}_d)$ for each $k \in \mathbb{K}$ and $i \in [n_k]$.

Based on Assumptions 1 and 2 and Proposition 1, the following strong duality results hold.

Theorem 1 (Strong Duality). *For the convex risk function Φ and Sinkhorn discrepancy $\mathcal{S}_\epsilon$, the following results hold:*

(I) *If Assumptions 1 and 2(I) hold, then*

$$\text{Optimal Value of (SDRHT)} = \max_{\mathbb{P}_0 \in \mathcal{B}_{\rho_0}(\widehat{\mathbb{P}}_0), \mathbb{P}_1 \in \mathcal{B}_{\rho_1}(\widehat{\mathbb{P}}_1)} \inf_{\phi \in \mathcal{F}} \Phi(\phi; \mathbb{P}_0, \mathbb{P}_1);$$

(II) *If Assumptions 1 and 2 hold, then*

$$\text{Optimal Value of (soft-SDRHT)} = \max_{\mathbb{P}_0, \mathbb{P}_1 \in \mathcal{P}(\Omega)} \inf_{\phi \in \mathcal{F}} \Phi(\phi; \mathbb{P}_0, \mathbb{P}_1) - \sum_{k \in \mathbb{K}} \lambda_k \mathcal{S}_\epsilon(\widehat{\mathbb{P}}_k, \mathbb{P}_k),$$

where $\lambda_k > 0$ for any $k \in \mathbb{K}$.

The proof of Theorem 1 is based on the lopsided minimax theorem [7, Theorem 5.15], and is provided in Appendix B.2. We prove this by showing that the Sinkhorn-based ambiguity sets are weakly compact and the expected loss is weakly upper semicontinuous.

Theorem 1 allows us to eliminate the detector minimization before optimizing over the least-favorable distributions. Thus, for generating function $\ell(t) = \exp(t)$, problem (soft-SDRHT) is equivalent to

$$\max_{\mathbb{P}_0, \mathbb{P}_1 \in \mathcal{P}(\Omega)} \Phi(\phi_{\text{opt}}; \mathbb{P}_0, \mathbb{P}_1) - \sum_{k \in \mathbb{K}} \lambda_k \mathcal{S}_\epsilon(\widehat{\mathbb{P}}_k, \mathbb{P}_k). \quad (5)$$

where

$$\Phi(\phi_{\text{opt}}; \mathbb{P}_0, \mathbb{P}_1) = \sum_{k \in \mathbb{K}} \mathbb{E}_{\omega \sim \mathbb{P}_k} \left[\exp \left[(-1)^{k+1} \cdot \frac{1}{2} \left(\log p_0(\omega) - \log p_1(\omega) \right) \right] \right]. \quad (6)$$

For a testing sample $\omega \in \Omega$, it suffices to evaluate the log-density of least-favorable distributions $\mathbb{P}_0$ and $\mathbb{P}_1$ at ω , and then make a decision by our explicitly optimal detector. Thus, the robust detector retains the likelihood-ratio form, but its densities are those of the least-favorable distributions. Problem (5) nevertheless remains infinite-dimensional because the least-favorable distributions are still optimization variables. In the next subsection, we characterize these distributions using parameterized convex-gradient transport maps.

2.2 Reformulation via Convex-Gradient Transport Maps

By setting the transport cost to the quadratic function $c(\mathbf{x}, \mathbf{z}) = \frac{1}{2}\|\mathbf{x} - \mathbf{z}\|_2^2$, we can leverage the following fundamental result from optimal transport [32] to simplify Problem (5).

Theorem 2 (Brenier’s Theorem). *Let μ, ν be two finite-second-moment probability measures on $\mathbb{R}^d$, and assume that μ is absolutely continuous with respect to the Lebesgue measure. Then there exists a proper convex function $\varphi : \mathbb{R}^d \rightarrow \mathbb{R}$ such that $(\nabla\varphi)_\# \mu = \nu$. In addition, $\nabla\varphi$ is the unique solution to*

$$\inf_{T: T_\# \mu = \nu} \mathbb{E}_\mu \|\mathbf{x} - T(\mathbf{x})\|_2^2,$$

and $\gamma_{\text{opt}} \sim (\text{Id}, \nabla\varphi)_\# \mu$ is the unique solution to

$$\inf_{\gamma \in \Gamma(\mu, \nu)} \mathbb{E}_\gamma \|\mathbf{x} - \mathbf{y}\|_2^2.$$

Theorem 2 shows that the gradient of a convex potential can characterize the transport map between continuous distributions. Proposition 1 and Theorem 2 enable us to reformulate (5) as an equivalent optimization problem over convex gradient transport maps, i.e.,

$$\max_{\varphi_0, \varphi_1 \in \mathcal{F}_{\text{CVX}}} \left\{ \Phi(\phi_{\text{opt}}; \mathbb{P}_0, \mathbb{P}_1) - \sum_{k \in \mathbb{K}} \lambda_k \mathcal{S}_\epsilon(\widehat{\mathbb{P}}_k, \mathbb{P}_k) : \mathbb{P}_k = (\nabla\varphi_k)_\# \mathbb{P}_k^\epsilon \right\}, \quad (7)$$

where $\mathcal{F}_{\text{CVX}}$ is the set of convex functions $\varphi : \mathbb{R}^d \rightarrow \mathbb{R} \cup \{+\infty\}$ and the distribution $\mathbb{P}_k^\epsilon$ for each $k \in \mathbb{K}$ is defined in Eq. (4). Note that we use continuous distributions $\mathbb{P}_k^\epsilon$ as source distributions of the transport maps for all $k \in \mathbb{K}$, thanks to the connection between Sinkhorn discrepancy and KL divergence shown in Proposition 1.

For computation, we restrict the convex function $\varphi_k \in C^2(\mathbb{R}^d)$ to be strongly convex, and $\nabla\varphi_k$ to be a global diffeomorphism for each $k \in \mathbb{K}$. In formulation (7), given strongly convex functions φ_0 and φ_1 , the log-density terms in risk function $\Phi(\phi_{\text{opt}}; \mathbb{P}_0, \mathbb{P}_1)$ can be explicitly computed by the log-density form of the change-of-variable formula [23]. Specifically, for any sample $\omega_k \sim \mathbb{P}_k$, we have

$$\log p_k(\omega_k) = \log p_k^\epsilon(z_k) - \log \left| \det \nabla^2 \varphi_k(z_k) \right|, \quad \forall k \in \mathbb{K}, \quad (8)$$

where $z_k = \nabla\varphi_k^{-1}(\omega_k)$, and p_k^ϵ represents the probability density function of $\mathbb{P}_k^\epsilon$ for each $k \in \mathbb{K}$. Based on the Fenchel–Young inequality, the gradient of a differentiable strongly convex function φ is the inverse of the gradient of its convex conjugate, i.e., $\nabla\varphi^{-1} = \nabla\varphi^*$. In this case, we have $z_k = \arg \max_{z \in \mathbb{R}^d} \{\omega_k^\top z - \varphi_k(z)\}$. Similarly, to evaluate the log-density of sample $\omega_k \sim \mathbb{P}_k$ at the other distribution $\mathbb{P}_{|1-k|}^\epsilon$, i.e., $\log p_{|1-k|}(\omega_k)$, we trace ω_k back to its corresponding source sample $\tilde{z}_k$ in $\mathbb{P}_{|1-k|}^\epsilon$:

$$\tilde{z}_k = \nabla\varphi_{|1-k|}^{-1}(\omega_k) = \arg \max_{z \in \mathbb{R}^d} \left\{ \omega_k^\top z - \varphi_{|1-k|}(z) \right\}, \quad \forall k \in \mathbb{K}, \quad (9)$$

and then we have

$$\log p_{|1-k|}(\omega_k) = \log p_{|1-k|}^\epsilon(\tilde{z}_k) - \log \left| \det \nabla^2 \varphi_{|1-k|}(\tilde{z}_k) \right|, \quad \forall k \in \mathbb{K}. \quad (10)$$

Since the potential φ_k in (7) is assumed to be twice differentiable and strongly convex, its Hessian determinant is strictly positive, i.e., $|\det \nabla^2 \varphi_k(z)| = \det \nabla^2 \varphi_k(z)$, and the log-determinant operators in Eqs. (8) and (10) are well defined. In high-dimensional applications, the exact values of the log-determinant-Hessian operators are computationally expensive, but can be estimated efficiently by the stochastic Lanczos quadrature (SLQ) in [33, 34].

Although formulation (7) exposes the transport structure of the least-favorable distributions, it remains an infinite-dimensional optimization. Moreover, the induced densities, inverse maps, and log-determinant terms are generally unavailable in closed form. However, this formulation offers a new perspective on generating least-favorable distributions by characterizing the gradients of convex potentials. In the next section, we propose a generative method to solve (7).

3 A HyCNN-Based Generative Method

Flow-based generative models represent a target distribution as the pushforward of a simple reference distribution through an invertible, differentiable map, allowing efficient density evaluation and sample generation [23]. In flow-based DRO, these models can also generate least-favorable distributions when guided by proper loss functions [26]. For problem (7), a natural approach is to parameterize the convex potentials by the Input Convex Neural Network (ICNN) [34, 35]. In this section, we develop a generative framework that parametrizes convex potentials by a variant of ICNN, i.e., *Hyper Input Convex Neural Network (HyCNN)*, extending the classical ICNN with more flexible architectures [36].

3.1 The Architecture of HyCNN

In this subsection, we first introduce the HyCNN architecture. Let the number of layers of HyCNN be L , and the output of HyCNN given an input $\mathbf{x} \in \mathbb{R}^d$ be $\text{HyCNN}(\mathbf{x}) \in \mathbb{R}$. Figure 2 illustrates the architecture of HyCNN, including the detailed structure of each block. The forward propagation process of HyCNN can be reformulated as follows

$$\text{HyCNN}(\mathbf{x}) := y_{L+1} = \mathbf{V}_L \mathbf{y}_L + \mathbf{W}_L \mathbf{x} + b_L, \quad (11a)$$

$$\mathbf{y}_{l+1} := \bar{\sigma} \left(\mathbf{V}_l^{(1)} \mathbf{y}_l + \mathbf{W}_l^{(1)} \mathbf{x} + \mathbf{b}_l^{(1)}, \mathbf{V}_l^{(2)} \mathbf{y}_l + \mathbf{W}_l^{(2)} \mathbf{x} + \mathbf{b}_l^{(2)} \right), \quad \forall l \in \{0, \dots, L-1\}. \quad (11b)$$

where $\{\mathbf{V}, \mathbf{W}, \mathbf{b}\}$ are trainable parameters, $\mathbf{y}_0 := \mathbf{0}$, and $\bar{\sigma} : \mathbb{R}^2 \rightarrow \mathbb{R}$ is a scalar activation function that is convex and component-wise non-decreasing, applied element-wise to vector inputs.

As shown in Figure 2, in each hidden layer $l \geq 1$ of HyCNN, a “passthrough” layer (red arrow) directly connects the input vector with the hidden units. Since the nonnegative sums of convex functions are also convex and the composition of a convex and convex non-decreasing function is also convex, according to Proposition 2.1 in [36], this construction ensures the convexity of HyCNNs.

Proposition 2 (Convexity of HyCNN [36]). *Suppose that $\bar{\sigma} : \mathbb{R}^2 \rightarrow \mathbb{R}$ is convex and component-wise non-decreasing, and weight matrices $\{\mathbf{V}_l^{(1)}, \mathbf{V}_l^{(2)}\}_{l=0}^{L-1}$ and $\mathbf{V}_L$ are component-wise nonnegative. Then, $\text{HyCNN}(\mathbf{x})$ in Eq. (11a) is convex in $\mathbf{x}$.*

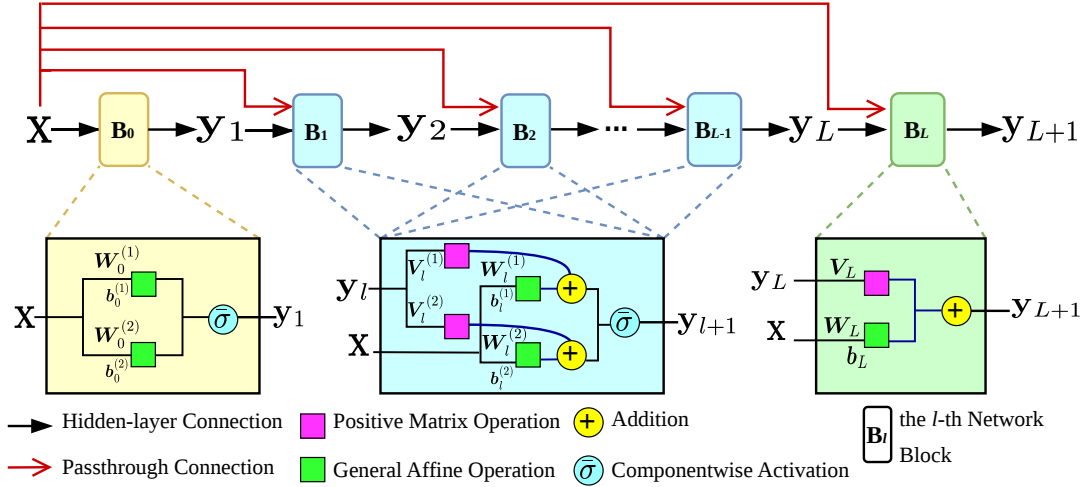


Figure 2: The architecture of HyCNN

In this paper, the non-negativity of weight matrices $\{V_l^{(1)}, V_l^{(2)}\}_{l=0}^{L-1}$ and V_L is enforced by soft-plus operators. Note that when we choose $\bar{\sigma}(a_1, a_2) = \sigma(a_1)$, where $\sigma : \mathbb{R} \rightarrow \mathbb{R}$ is an activation function that is convex and non-decreasing (e.g., ReLU), the HyCNN reduces to the standard ICNN architecture [35]. Thus, HyCNNs can be interpreted as a generalization of ICNNs. When we choose a maxout unit [37] as our activation function, i.e., $\bar{\sigma}(a_1, a_2) = \max(a_1, a_2)$, the resulting HyCNN is piecewise-affine but is nondifferentiable at ties. Therefore, to ensure the trained HyCNN is differentiable, we use a smooth log-sum-exp approximation of the maxout activation function with parameter $\tau > 0$:

$$\bar{\sigma}_\tau(a_1, a_2) := \tau \cdot \log \left(e^{a_1/\tau} + e^{a_2/\tau} \right), \quad (12)$$

which is referred to as the *smooth* maxout activation function below.

3.2 Approximation for Convex Potentials

In this subsection, we present the procedure for training HyCNNs. Given sample $z \in \Omega$, denote the two HyCNNs we use for the problem (7) by $\text{HyCNN}_0(z)$ and $\text{HyCNN}_1(z)$, respectively. Denote the approximated convex function by $\tilde{\varphi}_k$ for each $k \in \mathbb{K}$. To ensure the gradient of convex potentials is invertible, as the smooth maxout activation makes HyCNN_k twice continuously differentiable, we further add a positive quadratic term to ensure strong convexity, i.e., $\tilde{\varphi}_k(z) = \text{HyCNN}_k(z) + \frac{\alpha}{2} \|z\|_2^2$ where α is a sufficient small positive number, such that the Hessian matrix $H_{\tilde{\varphi}_k} \geq \alpha I > 0$. Unlike block-wise progressive methods that train a series of maps in [26], we directly learn one convex-gradient map for each hypothesis, since a composition of convex-gradient maps may not be the gradient of a convex potential.

The training process is given by the following Algorithm 1. For the gradient of the objective function, in Step 5, we propose two N -sample estimators to approximate the gradient of the risk

function and the Sinkhorn discrepancy, respectively; see Algorithms 3 and 4 in the next subsection for details. As shown, to generate the samples for estimators, we first draw N samples from the continuous reference distribution $\mathbb{P}_k^\epsilon$ (Step 2), and then map them to the target distribution space by the gradient of trained convex potentials (Step 3). In Step 6, the stochastic gradient ascent algorithm is used to update the network parameters.

Algorithm 1 Convex Potentials Approximation based on HyCNNs.

Require: Number of epochs T , training samples $\{\mathbf{x}_k^{(i)}\}_{i=1}^{n_k}$, reference distributions $\mathbb{P}_{k,0} = \mathbb{P}_k^\epsilon$.

Output: Convex potential approximations $\tilde{\varphi}_{k,T}$ for all $k \in \mathbb{K}$.

- 1: Initialize strongly convex potential $\tilde{\varphi}_{k,0}$ with parameters $\boldsymbol{\theta}_{k,0} \in \Theta$ for each $k \in \mathbb{K}$;
- 2: **for** $t = 0, \dots, T-1$ **do**
- 3: Draw N i.i.d. samples $\{\mathbf{z}_k^{(i)}\}_{i=1}^N$ from initial distribution $\mathbb{P}_{k,0}$ for each $k \in \mathbb{K}$;
- 4: Generate samples $\{\boldsymbol{\omega}_{k,t}^{(i)}\}_{i=1}^N \leftarrow \{\nabla \tilde{\varphi}_{k,t}(\mathbf{z}_k^{(i)})\}_{i=1}^N$ in target distribution $\mathbb{P}_{k,t}$ for each $k \in \mathbb{K}$;
- 5: Estimate the gradient of objective function of (7) by Algorithms 3 and 4:

$$\begin{aligned} \nabla_{\boldsymbol{\theta}_{k,t}} F(\boldsymbol{\theta}_{0,t}, \boldsymbol{\theta}_{1,t}) &\leftarrow \text{Risk_Gradient_Estimator}(\{\boldsymbol{\omega}_0^{(i)}\}_{i=1}^N, \{\boldsymbol{\omega}_1^{(i)}\}_{i=1}^N, \tilde{\varphi}_{0,t}, \tilde{\varphi}_{1,t}, \boldsymbol{\theta}_{k,t}) \\ &\quad - \nabla_{\boldsymbol{\theta}_{k,t}} \sum_{k \in \mathbb{K}} \lambda_k \cdot \text{Sinkhorn_Estimator}(\{\mathbf{x}_k^{(i)}\}_{i=1}^{n_k}, \{\boldsymbol{\omega}_{k,t}^{(i)}\}_{i=1}^N), \quad \forall k \in \mathbb{K}; \end{aligned}$$

- 6: Update $\boldsymbol{\theta}_{k,t+1} \leftarrow \boldsymbol{\theta}_{k,t} + \eta_t \cdot \nabla_{\boldsymbol{\theta}_{k,t}} F(\boldsymbol{\theta}_{0,t}, \boldsymbol{\theta}_{1,t})$, where η_t is the learning rate;
 - 7: **end for**
-

3.3 Stochastic Gradient Estimation

Next, we develop stochastic estimators required by Algorithm 1 to solve Problem (7). Let $\{\mathbf{z}_0^{(i)}\}_{i=1}^N$ and $\{\mathbf{z}_1^{(i)}\}_{i=1}^N$ be N independent and identically distributed (i.i.d.) samples from reference distributions $\mathbb{P}_0^\epsilon$ and $\mathbb{P}_1^\epsilon$ (see Eq. 4), respectively. These points are subsequently transformed into samples from the target distributions $\mathbb{P}_0$ and $\mathbb{P}_1$ via transport maps, i.e., $\boldsymbol{\omega}_0 = \nabla \varphi_0(\mathbf{z}_0)$ and $\boldsymbol{\omega}_1 = \nabla \varphi_1(\mathbf{z}_1)$, where the underlying distributions satisfy $\mathbb{P}_0 = [\nabla \varphi_0]_{\#} \mathbb{P}_0^\epsilon$ and $\mathbb{P}_1 = [\nabla \varphi_1]_{\#} \mathbb{P}_1^\epsilon$.

3.3.1 Risk Gradient Estimator

In this part, we present the gradient estimator of the risk function with respect to the parameters of HyCNNs in Algorithm 1. For every $k \in \mathbb{K}$ and $\boldsymbol{\omega}_k \sim \mathbb{P}_k$, we define the basic element of the risk function as

$$h_k(\boldsymbol{\omega}_k, \boldsymbol{\theta}_0, \boldsymbol{\theta}_1) := \exp \left[\frac{1}{2} (\log p_{|1-k|}(\boldsymbol{\omega}_k) - \log p_k(\boldsymbol{\omega}_k)) \right], \quad \forall k \in \mathbb{K}.$$

As an example, given convex functions $\tilde{\varphi}_0$ and $\tilde{\varphi}_1$ with parameters $\boldsymbol{\theta}_0$ and $\boldsymbol{\theta}_1$, respectively, the gradients of $h_0(\boldsymbol{\omega}_0, \boldsymbol{\theta}_0, \boldsymbol{\theta}_1)$ with respect to $\boldsymbol{\theta}_0$ and $\boldsymbol{\theta}_1$ are given by

$$\frac{\partial h_0(\boldsymbol{\omega}_0, \boldsymbol{\theta}_0, \boldsymbol{\theta}_1)}{\partial \boldsymbol{\theta}_0} = \frac{h_0(\boldsymbol{\omega}_0, \boldsymbol{\theta}_0, \boldsymbol{\theta}_1)}{2} \cdot \left[\left(\frac{\partial \boldsymbol{\omega}_0}{\partial \boldsymbol{\theta}_0} \right)^\top [\nabla^2 \tilde{\varphi}_1(\tilde{\mathbf{z}}_0)]^{-1} \mathbf{g} + \frac{\partial}{\partial \boldsymbol{\theta}_0} \log \det \nabla^2 \tilde{\varphi}_0(\mathbf{z}_0) \right], \quad (13a)$$

$$\frac{\partial h_0(\boldsymbol{\omega}_0, \boldsymbol{\theta}_0, \boldsymbol{\theta}_1)}{\partial \boldsymbol{\theta}_1} = -\frac{h_0(\boldsymbol{\omega}_0, \boldsymbol{\theta}_0, \boldsymbol{\theta}_1)}{2} \cdot \left[\left(\frac{\partial \nabla \tilde{\varphi}_1(\tilde{\mathbf{z}}_0)}{\partial \boldsymbol{\theta}_1} \right)^\top [\nabla^2 \tilde{\varphi}_1(\tilde{\mathbf{z}}_0)]^{-1} \mathbf{g} + \frac{\partial}{\partial \boldsymbol{\theta}_1} \log \det \nabla^2 \tilde{\varphi}_1(\tilde{\mathbf{z}}_0) \right], \quad (13b)$$

where $\mathbf{g} := \frac{\partial \log p_1(\omega_0)}{\partial \tilde{\mathbf{z}}_0} = \frac{\partial [\log p_1^e(\tilde{\mathbf{z}}_0) - \log \det \nabla^2 \tilde{\varphi}_1(\tilde{\mathbf{z}}_0)]}{\partial \tilde{\mathbf{z}}_0}$, and $\tilde{\mathbf{z}}_k$ is defined in Eq. (9) for each $k \in \mathbb{K}$. For brevity, the detailed gradient derivation can be found in Appendix C.1. Eqs. (13a) and (13a) indicate that directly differentiating the risk requires backpropagation through inverse maps and the log-determinant-Hessian (LDH), which is computationally expensive. Therefore, in the following, we develop a matrix-free stochastic gradient estimator for the LDH in Algorithm 2, and then present the estimator for the risk function in Algorithm 3.

For the gradient of LDH, we first introduce an $\mathcal{O}(1)$ -memory estimator for the gradient of the log-determinant [34]. Specifically, for any invertible matrix H with parameters θ :

$$\frac{\partial}{\partial \theta} \log \det H = \text{tr}(H^{-1} \frac{\partial H}{\partial \theta}) = \mathbb{E}_{\mathbf{v}}[\mathbf{v}^\top H^{-1} \frac{\partial H}{\partial \theta} \mathbf{v}], \quad (14)$$

where the last equality uses the Hutchinson trace estimator [38] with an isotropic random vector (e.g., Rademacher or standard Gaussian) $\mathbf{v} \in \mathbb{R}^d$ satisfying $\mathbb{E}[\mathbf{v}] = 0$ and $\mathbb{E}[\mathbf{v}\mathbf{v}^\top] = \mathbf{I}$. With exact linear solves, the Hutchinson construction offers an unbiased estimator of the gradient of the log-determinant. In this problem, H is a symmetric positive definite Hessian matrix. Thus, $\mathbf{v}^\top H^{-1}$ in Eq. (14) can be efficiently solved by the following quadratic optimization problem

$$\min_{\mathbf{u}} \frac{1}{2} \mathbf{u}^\top H \mathbf{u} - \mathbf{v}^\top \mathbf{u}, \quad (15)$$

which has the unique minimizer $\mathbf{u}_{\text{opt}} = H^{-1} \mathbf{v} = (\mathbf{v}^\top H^{-1})^\top$. Then, the gradient of the LDH in our problem can be estimated as $\mathbb{E}_{\mathbf{v}}[\frac{\partial}{\partial \theta}(\mathbf{u}_{\text{opt}}^\top H \mathbf{v})]$. In this paper, we summarize this procedure as Algorithm 2.

Algorithm 2 GLDH_Estimator: Estimator for the Gradient of LDH.

Require: Symmetric positive definite matrix H , parameter θ , and sample size M .

- 1: Initialize estimator $\mathbf{E} = \mathbf{0}$ with the same dimension as θ ;
 - 2: **for** $m = 1, \dots, M$ **do**
 - 3: Sample Rademacher random vector $\mathbf{v}^{(m)}$ with d dimensions;
 - 4: $\mathbf{u}^{(m)} \leftarrow \text{ConjugateGradient}(H, \mathbf{v}^{(m)})$;
 - 5: $\bar{\mathbf{u}}^{(m)} \leftarrow \text{StopGradient}(\mathbf{u}^{(m)})$;
 - 6: $\mathbf{E} \leftarrow \mathbf{E} + \nabla_{\theta} \left[(\bar{\mathbf{u}}^{(m)})^\top H \mathbf{v}^{(m)} \right]$;
 - 7: **end for**
 - 8: **Return** $\mathbf{E}/M$.
-

In Step 4 of Algorithm 2, the function $\text{ConjugateGradient}(H, \mathbf{v}) := \arg \min_{\mathbf{u}} \frac{1}{2} \mathbf{u}^\top H \mathbf{u} - \mathbf{v}^\top \mathbf{u}$ represents that we use the conjugate gradient method to solve the unconstrained optimization problem (15). In Step 5, $\text{StopGradient}(\cdot)$ treats its argument as a constant during backpropagation. Thus, in Step 6, only $H(\theta)$ is differentiated, and the term $\nabla_{\theta}[(\bar{\mathbf{u}}^{(m)})^\top H \mathbf{v}^{(m)}]$ can be efficiently computed by Hessian-vector and vector-Jacobian products.

Then, the gradient of the risk function (6) can be estimated by Algorithm 3. Both functions ConjugateGradient in Step 5 and GLDH_Estimator have been defined in Algorithm 2. In Step 6, we compute the value of the log-determinant in $h_k(\omega_k, \theta_0, \theta_1)$ exactly when $d \leq 32$ and estimate it by the stochastic Lanczos quadrature (SLQ) [33, 34] when $d > 32$.

Algorithm 3 Risk_Gradient_Estimator: Gradient Estimator for the Risk Function.

Require: Generated samples $\{\omega_k^{(i)}\}_{i=1}^N$ and convex potentials $\tilde{\varphi}_k$ for all $k \in \mathbb{K}$, parameter θ .

- 1: Initialize estimator $E = \mathbf{0}$ with the same dimension as θ ;
- 2: **for** $k = 0, 1$ **do**
- 3: **for** $i = 1, \dots, N$ **do**
- 4: Compute $z_k^{(i)} = \nabla \tilde{\varphi}_k^{-1}(\omega_k^{(i)})$ and $\tilde{z}_k^{(i)} = \nabla \tilde{\varphi}_{|1-k|}^{-1}(\omega_k^{(i)})$ by Eq. (9);
- 5: Compute $u_k^{(i)} \leftarrow \text{ConjugateGradient}\left(\nabla^2 \tilde{\varphi}_{|1-k|}(\tilde{z}_k^{(i)}), \frac{\partial \log p_{|1-k|}(\omega_k^{(i)})}{\partial \tilde{z}_k^{(i)}}\right)$;
- 6: Update the gradient estimator

$$E \leftarrow E + \frac{h_k(\omega_k^{(i)}, \theta_0, \theta_1)}{2} \cdot \left\{ \left[\frac{\partial \left(\omega_k^{(i)} - \nabla \tilde{\varphi}_{|1-k|}(\tilde{z}_k^{(i)}) \right)}{\partial \theta} \right]^\top u_k^{(i)} \right. \\ \left. + \text{GLDH_Estimator}\left(\nabla^2 \tilde{\varphi}_k(z_k^{(i)}), \theta\right) - \text{GLDH_Estimator}\left(\nabla^2 \tilde{\varphi}_{|1-k|}(\tilde{z}_k^{(i)}), \theta\right) \right\};$$

- 7: **end for**
 - 8: **end for**
 - 9: **Return** E/N .
-

3.3.2 Sinkhorn Discrepancy Estimator

Computing the Sinkhorn discrepancy term $S_\epsilon(\widehat{\mathbb{P}}_k, \mathbb{P}_k)$ exactly (see Definition 2) in the objective function is challenging. In this part, we present a differentiable finite-sample estimator for the Sinkhorn discrepancy. The gradient of this estimator can be computed efficiently by backpropagation and is used in Algorithm 1.

We first introduce the following proposition regarding the dual formulation of Sinkhorn discrepancy [32, 39].

Proposition 3. *Let $\mathbb{P}$ and $\mathbb{Q}$ be probability measures with finite second moments. The strong dual formulation of Sinkhorn discrepancy $S_\epsilon(\mathbb{P}, \mathbb{Q})$ is given by*

$$\sup_{f \in L^1(\mathbb{P}), g \in L^1(\mathbb{Q})} \left\{ \int f(x) d\mathbb{P}(x) + \int g(z) d\mathbb{Q}(z) - \epsilon \iint \left[\exp\left(\frac{f(x) + g(z) - c(x, z)}{\epsilon}\right) - 1 \right] d\mathbb{P}(x) d\mathbb{Q}(z) \right\},$$

Let f_ϵ and g_ϵ be the solutions of the dual formulation, which is also known as optimal entropic potentials. Then the transport plan

$$\gamma_\epsilon(dx, dz) = \exp\left(\frac{f_\epsilon(x) + g_\epsilon(z) - c(x, z)}{\epsilon}\right) d\mathbb{P}(x) d\mathbb{Q}(z)$$

is the unique solution to the Sinkhorn discrepancy, and the optimal entropic potentials satisfy

$$\int \exp\left(\frac{f_\epsilon(x) + g_\epsilon(z) - c(x, z)}{\epsilon}\right) d\mathbb{P}(x) = 1, \quad \forall z \in \Omega, \quad (16a)$$

$$\int \exp\left(\frac{f_\epsilon(\mathbf{x}) + g_\epsilon(\mathbf{z}) - c(\mathbf{x}, \mathbf{z})}{\epsilon}\right) dQ(\mathbf{z}) = 1, \quad \forall \mathbf{x} \in \Omega, \quad (16b)$$

for Q -almost every $\mathbf{z}$ and $\mathbb{P}$ -almost every $\mathbf{x}$, respectively.

Since the target continuous distribution $\mathbb{P}_k$ is computationally intractable, we approximate it as empirical distributions formed by the generated samples, denoted by $\widetilde{\mathbb{P}}_0 = \frac{1}{N} \sum_{i=1}^N \delta_{\omega_0^{(i)}}$ and $\widetilde{\mathbb{P}}_1 = \frac{1}{N} \sum_{i=1}^N \delta_{\omega_1^{(i)}}$. Therefore, it suffices to compute the value $\mathcal{S}_{\epsilon, (n_k, N)}(\widetilde{\mathbb{P}}_k, \widetilde{\mathbb{P}}_k)$. The optimal entropic potentials between discrete distributions $\widehat{\mathbb{P}}_k$ and $\widetilde{\mathbb{P}}_k$, denoted by $f_{\epsilon, k, (n_k, N)}$ and $g_{\epsilon, k, (n_k, N)}$, can be obtained by Sinkhorn's algorithm [20, 39], leading to the following estimator for the barycentric projection of the optimal entropic coupling:

$$T_{\epsilon, k, (n_k, N)}(\mathbf{x}) = \frac{\frac{1}{N} \sum_{i=1}^N \omega_k^{(i)} e^{\left(\frac{g_{\epsilon, k, (n_k, N)}(\omega_k^{(i)}) - c(\mathbf{x}, \omega_k^{(i)})}{\epsilon}\right)}}{\frac{1}{N} \sum_{i=1}^N e^{\left(\frac{g_{\epsilon, k, (n_k, N)}(\omega_k^{(i)}) - c(\mathbf{x}, \omega_k^{(i)})}{\epsilon}\right)}}, \quad \forall k \in \mathbb{K}.$$

However, the discrete dual potential $g_{\epsilon, k, (n_k, N)}$ is initially defined only on the generated support points and is thus not differentiable for all samples in Ω . To make the optimized empirical Sinkhorn discrepancy differentiable with respect to the generated sample locations, we extend g to the general continuous sample space for fixed f :

$$g_{\epsilon, k}^c(\omega) = -\epsilon \log \left(\frac{1}{n_k} \sum_{j=1}^{n_k} \exp \left(\frac{f_k(\mathbf{x}_k^{(j)}) - c(\mathbf{x}_k^{(j)}, \omega)}{\epsilon} \right) \right), \quad \forall \omega \in \Omega. \quad (17)$$

This extension agrees with the original potential g on the empirical support at convergence while making it differentiable. Therefore, our differentiable estimator for the Sinkhorn discrepancy is given by

$$\widehat{\mathcal{S}}_{\epsilon, (n_k, N)}(\widehat{\mathbb{P}}_k, \mathbb{P}_k) := \mathcal{S}_{\epsilon, (n_k, N)}(\widehat{\mathbb{P}}_k, \widetilde{\mathbb{P}}_k) = \frac{1}{n_k} \sum_{j=1}^{n_k} f_{\epsilon, k, (n_k, N)}(\mathbf{x}_k^{(j)}) + \frac{1}{N} \sum_{i=1}^N g_{\epsilon, k}^c(\omega_k^{(i)}), \quad (18)$$

and its gradient can be computed by backpropagation (see Appendix C.2).

Therefore, the pseudo-code for estimating the differentiable Sinkhorn discrepancy is shown in Algorithm 4. Steps 1-6 in Algorithm 4 are essentially Sinkhorn's algorithm [39], and thus the computational complexity of each iteration is $\mathcal{O}(n_k \cdot N)$. This alternating algorithm computes the entropic potentials by iteratively updating f and g while satisfying one of the two optimal marginal conditions in Eq. (16).

4 Representation Power of HyCNNs

This section provides theoretical analyses of the representation capacity of HyCNNs. Specifically, we focus on their representation power for both convex functions and their gradients, as well as the distributional universality of their induced transport maps.

Algorithm 4 Sinkhorn_Estimator: Estimator for the Sinkhorn Discrepancy.

Require: Training samples $\{\mathbf{x}_k^{(i)}\}_{i=1}^{n_k}$ from $\widehat{\mathbb{P}}_k$ and generated samples $\{\boldsymbol{\omega}_k^{(i)}\}_{i=1}^N$.

- 1: Initialize $i = 0$ and $f_{\epsilon,k}^{(0)} = 0$ for each k ;
 - 2: **while** not converged **do**
 - 3: update $g_{\epsilon,k}^{(i+1)}(\boldsymbol{\omega}_k) \leftarrow -\epsilon \log \frac{1}{n_k} \sum_{j=1}^{n_k} \exp\left(\frac{f_{\epsilon,k}^{(i)}(\mathbf{x}_k^{(j)}) - c(\mathbf{x}_k^{(j)}, \boldsymbol{\omega}_k)}{\epsilon}\right)$ for each $k \in \mathbb{K}$;
 - 4: update $f_{\epsilon,k}^{(i+1)}(\mathbf{x}_k) \leftarrow -\epsilon \log \frac{1}{N} \sum_{j=1}^N \exp\left(\frac{g_{\epsilon,k}^{(i+1)}(\boldsymbol{\omega}_k^{(j)}) - c(\mathbf{x}_k, \boldsymbol{\omega}_k^{(j)})}{\epsilon}\right)$ for each $k \in \mathbb{K}$;
 - 5: $i \leftarrow i + 1$;
 - 6: **end while**
 - 7: Extend the entropic potential $g_{\epsilon,k}^{(i+1)}(\boldsymbol{\omega}_k)$ to continuous space by Eq. (17);
 - 8: **Return** $\widetilde{S}_{\epsilon, (n_k, N)}(\widehat{\mathbb{P}}_k, \mathbb{P}_{k,t})$ computed by Eq. (18).
-

We refer to HyCNNs equipped with *smooth* maxout activation functions $\bar{\sigma}_\tau$ (see Eq. (12)) as *smooth HyCNNs*. We first show the representation power of both standard and smooth HyCNNs. Let L and h denote the number of layers and the width at each hidden layer of HyCNNs, respectively. Since we take the strongly convex coefficient $\alpha \rightarrow 0$, these results hold in our strongly convex setting.

In the following results, we take the strongly convex coefficient $\alpha \rightarrow 0$.

Theorem 3 (Representation Power of (Smooth) HyCNNs). *For any L_f -Lipschitz convex function $f : \mathbb{R}^d \rightarrow \mathbb{R}$, compact convex set $\mathbb{D} \subset \mathbb{R}^d$, and $\epsilon > 0$, assume that $\mathbb{D}$ has a diameter $D \in (0, \infty)$ such that $\|\mathbf{x}_1 - \mathbf{x}_2\| \leq D$ for any $\mathbf{x}_1, \mathbf{x}_2 \in \mathbb{D}$, and then:*

- (I) *There exists a hyper input convex neural network $\widehat{\text{HyCNN}}$ with maxout activation functions satisfying $\sup_{\mathbf{x} \in \mathbb{D}} |\widehat{\text{HyCNN}}(\mathbf{x}) - f(\mathbf{x})| \leq \epsilon$ and $L \cdot h = \mathcal{O}((\frac{DL_f}{\epsilon})^d)$.*
- (II) *There exists a smooth HyCNN with a sufficiently small parameter $\tau > 0$, denoted by $\widehat{\text{HyCNN}}^\tau$, satisfying $\sup_{\mathbf{x} \in \mathbb{D}} |\widehat{\text{HyCNN}}^\tau(\mathbf{x}) - f(\mathbf{x})| \leq \epsilon$ and $L \cdot h = \mathcal{O}((\frac{DL_f}{\epsilon})^d)$.*
- (III) *There exists a sequence of smooth HyCNNs, denoted by $\{\widehat{\text{HyCNN}}_n^\tau\}_{n=1}^\infty$, satisfying $\widehat{\text{HyCNN}}_n^\tau \rightarrow f$ uniformly on $\mathbb{D}$, when τ is sufficiently small and $L \cdot h = \mathcal{O}((\frac{DL_f}{\epsilon})^d)$.*

The proof of Theorem 3 is provided in Appendix B.3. Compared to existing qualitative representation power theorems of ICNNs [34, 40], Theorem 3 extends this result to HyCNNs and provides quantitative guidelines for hyperparameter selection.

Although these properties of (smooth) HyCNNs show that they can uniformly approximate convex potentials on a compact set, convergence of potentials does not generally imply convergence of the gradient fields [34]. Theorem 4 shows that the smooth HyCNNs approximate the gradient of convex functions well.

Theorem 4 (Gradient Convergence of Smooth HyCNNs). *Let $f : \mathbb{R}^d \rightarrow \mathbb{R}$ be a Lipschitz continuous convex function, and $\{\widehat{\text{HyCNN}}_n^\tau\}_{n=1}^\infty$ be a sequence of smooth HyCNNs with hyperparameters $L_n \cdot h_n = \mathcal{O}(\epsilon_n^{-d})$ for any $\epsilon_n > 0$, where $\widehat{\text{HyCNN}}_n^\tau : \mathbb{R}^d \rightarrow \mathbb{R}$. Assume that $\widehat{\text{HyCNN}}_n^\tau \rightarrow f$. Then*

- (I) For almost every $\mathbf{x} \in \mathbb{R}^d$, $\nabla \widehat{\text{HyCNN}}_n^\tau(\mathbf{x}) \rightarrow \nabla f(\mathbf{x})$.
- (II) If f is Lipschitz smooth, the gradient sequence $\nabla \widehat{\text{HyCNN}}_n$ converges to ∇f uniformly at a rate of $\mathcal{O}(\sqrt{\varepsilon_n})$.

The proof of Theorem 4 is provided in Appendix B.4. Generally, Theorem 4 holds for any differentiable convex function sequence. Note that the uniform convergence assumption in Theorem 4 can be verified using the representation power of smooth HyCNNs on compact sets (see Theorem 3(III)).

Based on the results above and Brenier’s theorem, the following Theorem 5 shows that smooth HyCNNs are distributionally universal in our problem.

Theorem 5 (Distributional Universality of Smooth HyCNNs). *Let $\mu, \nu \in \mathcal{P}(\Omega)$ be regarded as probability measures on $\mathbb{R}^d$ supported on Ω . Assume that μ is absolutely continuous with respect to the d -dimensional Lebesgue measure on $\mathbb{R}^d$, and that μ and ν have finite second moments. Then there exists a sequence of smooth HyCNNs, denoted by $\{\widehat{\text{HyCNN}}_n^\tau\}_{n=1}^\infty$, such that $(\nabla \widehat{\text{HyCNN}}_n^\tau)_\# \mu \Rightarrow \nu$, where $\Rightarrow$ denotes weak convergence of probability measures on $\mathbb{R}^d$.*

The proof of Theorem 5 is provided in Appendix B.5. It demonstrates that smooth HyCNN gradient maps are distributionally universal for approximating target distributions. For a reference distribution $\mathbb{P}^\epsilon \in \mathcal{P}(\Omega)$ and any candidate least-favorable distribution $\mathbb{P}^\star \in \mathcal{P}(\Omega)$ in our problem, assume that both $\mathbb{P}^\epsilon$ and $\mathbb{P}^\star$ have finite second moments, and $\mathbb{P}^\epsilon$ is absolutely continuous with respect to the Lebesgue measure on $\mathbb{R}^d$. Then, the gradient map of smooth HyCNNs is distributionally universal for approximating any candidate least-favorable distribution.

5 Experiments

In this section, we conduct experiments to assess the effectiveness of the proposed method on datasets of varying training sample sizes and dimensions. Specifically, we evaluate its performance on small sample sizes using a synthetic Gaussian mixture dataset (Section 5.1) and the MNIST dataset (Section 5.2), and on large sample sizes using the Higgs dataset (Section 5.3). We also provide a qualitative illustration of our method on a high-dimensional human face generation task (Section 5.4). All experiments were conducted on a laptop with an NVIDIA GeForce RTX 4060 GPU using Python. Compared to the existing methods, our approach can handle hypothesis testing problems across different sample sizes and dimensions, achieving higher accuracy and robustness. Codes and experimental results are available at <https://github.com/Arthas819/SDRHT>.

5.1 Synthetic Gaussian Mixture Dataset

In this subsection, we consider a synthetic Gaussian mixture dataset based on [17]. Let the samples under the two hypotheses be generated from distributions $0.5\mathcal{N}(0.4\mathbf{e}, I_d) + 0.5\mathcal{N}(-0.4\mathbf{e}, I_d)$ and $0.5\mathcal{N}(0.4\mathbf{f}, I_d) + 0.5\mathcal{N}(-0.4\mathbf{f}, I_d)$, respectively. Here $\mathbf{e} \in \mathbb{R}^d$ is a vector with all entries equal to 1, and $\mathbf{f} \in \mathbb{R}^d$ is a vector with the first $\lceil d/2 \rceil$ entries equal to 1 and the remaining entries equal to -1 . Consider a setting with a small number of training samples $n_0 = n_1 = 10$, and then test

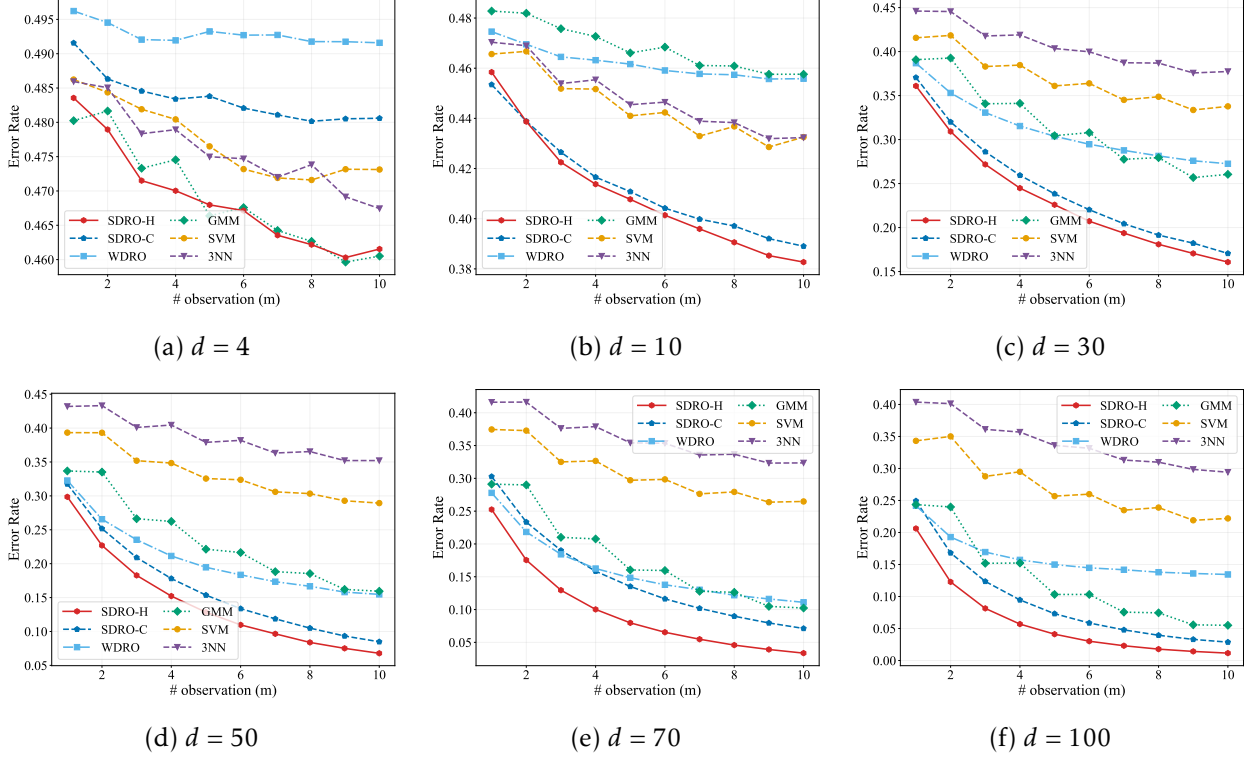


Figure 3: Error rate comparison on GMM dataset, averaged over 100 trials. **SDRO-H**: SDRHT method with HyCNN; **SDRO-C**: SDRHT method with standard ICNN.

on 1000 new samples from each mixture model. We compare our method against Wasserstein DRO-based hypothesis testing (WDRO) [17], Gaussian Mixture Model (GMM), kernel support vector machine (SVM) with radial basis function (RBF) kernel, and a three-layer neural network (3NN) with hidden layer dimensions (32,32,32). The details of the WDRO method are provided in Appendix D. All hyperparameters, including λ and ϵ in SDRO-H (SDRHT with HyCNN) and SDRO-C (SDRHT with standard ICNN), ambiguity set radii, and kernel bandwidth in WDRO, are determined by cross-validation.

Table 2: Computation time (s) of methods on Gaussian mixture dataset, averaged over 100 trials.

d	Computation Time (s)					
	SDRO-H	SDRO-C	WDRO	GMM	SVM	3NN
4	1.906	1.749	0.283	0.075	<0.000	0.052
10	1.507	1.802	0.275	0.002	0.001	0.040
30	1.292	1.900	0.271	0.002	<0.000	0.034
50	1.753	1.897	0.285	0.002	0.001	0.033
70	1.614	1.907	0.290	0.002	<0.000	0.026
100	1.725	1.844	0.327	0.002	<0.000	0.028

Figure 3 compares the classification error rate of selected methods over varying dimensions $d \in \{4, 10, 30, 50, 70, 100\}$. The horizontal axis in Figure 3 represents that there are m independent observations in a testing sample, and the decision is made by majority rule [17]. As illustrated, Hy-CNN is more suitable for approximating convex potentials in SDRO than standard ICNN (though the performance of SDRO-C remains competitive in Figures 3b-3f), and the SDRO-H method achieves the lowest classification error rate in almost all dimensions. These results demonstrate the superior performance of our method on hypothesis-testing tasks with small sample sizes compared to selected baselines. Table 2 shows the average computation time of all methods. Although the training time of SDRO is higher than that of other baselines, it remains under 2s and remains stable as d increases.

Figure 4 reports the accuracy of SDRO-H and WDRO across different hyperparameter combinations when $n_0 = n_1 = 10$ and $d = 100$. As illustrated, the accuracy under the best parameter combination of WDRO is inferior to that of SDRO-H. The result of SDRO-H indicates that one may improve accuracy by moderately increasing the penalty parameter λ and decreasing the entropic regularization parameter ϵ . This is because a small λ leads to excessive conservatism, and a large ϵ induces stronger smoothing and may dilute the correlation between samples and classification labels. The result of WDRO demonstrates that under an appropriate bandwidth selection, the classification accuracy is not sensitive to λ .

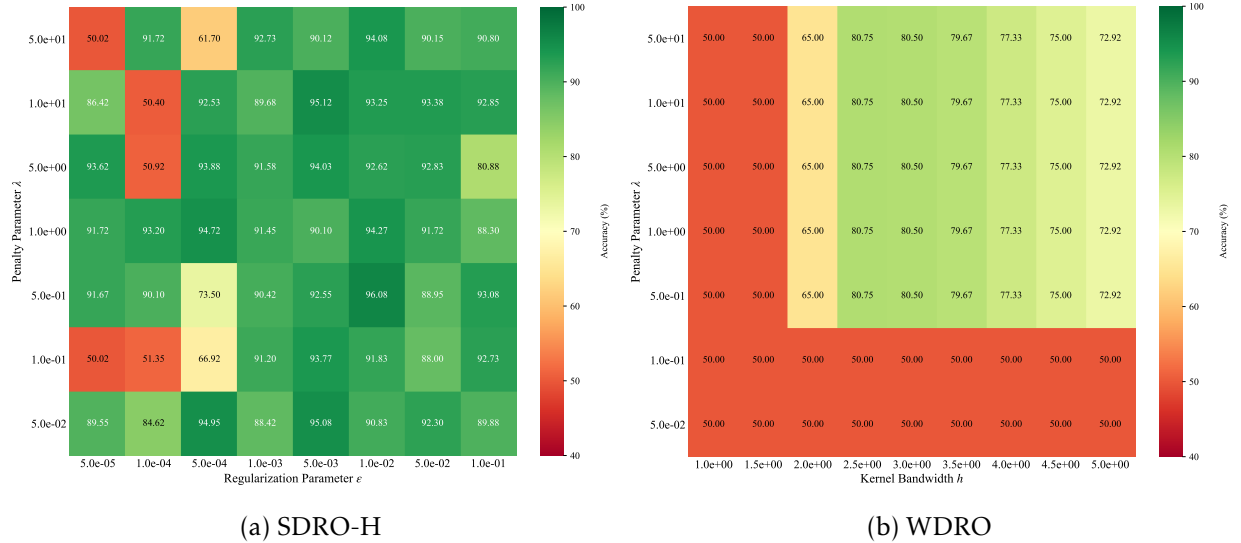


Figure 4: Average accuracy (%) of SDRO-H and WDRO on the GMM dataset with different hyperparameters ($n_0 = n_1 = 10, d = 100$).

To evaluate the adversarial robustness of our model on adversarial examples, we apply adversarial perturbations to the test data using the projected gradient descent (PGD) attack [26, 41]. Specifically, for a testing sample $\omega \sim \mathbb{P}_k$, the attacked sample for fixed detector ϕ_{opt} , defined by a point-wise perturbation problem, is given by

$$\omega_{\text{attacked}} := \omega + \delta_{\omega}^*, \quad \delta_{\omega}^* := \arg \max_{\|\delta_{\omega}\|_p \leq \Delta} L_k(\phi_{\text{opt}}, \omega + \delta_{\omega}). \quad (19)$$

Table 3: Robustness comparison on the GMM dataset after PGD attack, averaged over $m \in [10]$ and 5 trials (accuracy %, $n_0 = n_1 = 10, d = 50$). Bold represents the highest accuracy in each column.

Method	PGD- L_2 Attack				PGD- L_∞ Attack			
	$\Delta = 0.025$	$\Delta = 0.05$	$\Delta = 0.075$	$\Delta = 0.1$	$\Delta = 0.05$	$\Delta = 0.1$	$\Delta = 0.15$	$\Delta = 0.2$
SDRO-H	85.801	85.769	85.736	85.737	85.467	85.165	84.906	84.811
SDRO-C	82.333	82.347	82.332	82.309	82.175	82.018	81.867	81.774
WDRO	69.693	68.902	68.108	67.924	62.440	54.711	48.096	44.256
GMM	73.401	72.392	71.367	71.121	63.100	51.579	41.804	36.414
SVM	64.477	63.628	62.774	62.561	55.957	47.171	40.166	36.166
3NN	58.187	57.021	55.906	55.638	47.635	37.184	29.483	25.509

In the following, we assess robustness using L_2 and L_∞ norm constraints, i.e., $p = 2$ and $p \rightarrow \infty$ in Eq. (19). In Table 3, we compare the classification accuracy of all methods across different levels of Δ . As shown in Table 3, the accuracies of the proposed methods remain higher and more stable than those of the baselines as the attack intensity Δ increases under both L_2 and L_∞ constraints, and SDRO-H consistently yields the best test accuracy across all values of Δ .

5.2 MNIST Handwritten Digits Classification

This subsection considers the MNIST handwritten digits dataset ($d=784$). In each trial, we randomly select 5 digits as class 0, and the remaining as class 1. We select n' samples for each digit of each class to form the two training sets, and therefore $n_0 = n_1 = 5n'$. The testing data includes 1000 images from both classes. For this dataset, we further include logistic regression (LR) and Flow-based DRO (FDRO) [26] as baselines. The FDRO method, which characterizes least-favorable distributions from a generative perspective, is suitable for high-dimensional robust hypothesis testing and requires training convolutional autoencoders. In the following, the label ‘SDRO’ represents our SDRHT method with HyCNN. All hyperparameters are determined by cross-validation.

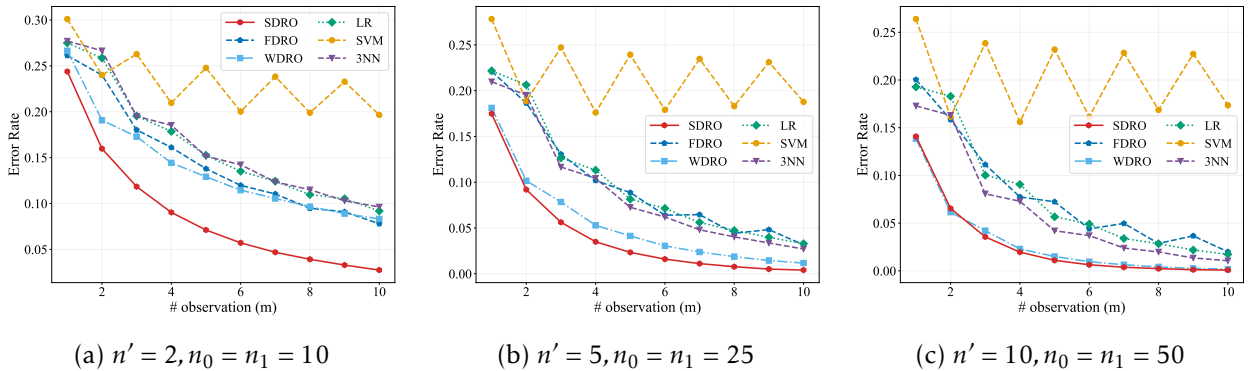


Figure 5: Error rate comparison on the MNIST dataset, averaged over 100 trials.

Figure 5 illustrates that the proposed method achieves a lower error rate than existing well-tuned DRO methods and common classifiers across different training set sizes. Table 4 shows the average computation time of all methods, illustrating that the training time of SDRO is acceptable and stable as n' increases. To evaluate robustness, we compare all methods under both L_2 and L_∞ PGD attacks (see Eq. (19)) in Table 5. As shown, SDRO achieves the highest accuracy across all non-zero attack intensities under both norms, indicating the superior robustness of our method. Under the PGD- L_∞ attack, the three DRO-based models demonstrate better robustness than other baselines.

Table 4: Computation time (s) of methods on the MNIST dataset, averaged over 100 trials.

n'	Computation Time (s)					
	SDRO	FDRO	WDRO	LR	SVM	3NN
2	1.943	4.821	0.022	0.008	<0.001	0.167
5	1.477	8.999	0.045	0.011	0.001	0.231
10	1.465	17.153	0.135	0.018	0.001	0.269

Table 5: Robustness comparison on MNIST dataset after PGD attack, averaged over $m \in [10]$ and 5 trials (accuracy %, $n_0 = n_1 = 25$).

Method	PGD- L_2 Attack				PGD- L_∞ Attack			
	$\Delta = 0.025$	$\Delta = 0.05$	$\Delta = 0.075$	$\Delta = 0.1$	$\Delta = 0.05$	$\Delta = 0.1$	$\Delta = 0.15$	$\Delta = 0.2$
SDRO	95.75	95.70	95.67	95.67	93.53	90.60	87.46	85.47
FDRO	91.04	90.93	90.87	90.87	86.51	80.92	75.28	71.91
WDRO	94.54	94.47	94.42	94.42	91.18	85.63	79.42	75.56
LR	91.15	90.98	90.89	90.89	84.79	75.80	67.27	62.34
SVM	76.67	76.50	76.44	76.44	72.65	69.02	66.21	64.73
3NN	91.39	91.21	91.11	91.11	82.10	68.34	55.00	48.09

We next visualize how the trained convex functions induce adversarial distributional attacks on their opposite classes. In each trial, we select two digits that are easily confused in handwritten form and place them in separate classes. To generate a more intuitive and clear trajectory, in the training phase, we first train a variational autoencoder, including two parts, Encoder and Decoder, to embed all images into a 32-dimensional latent space. Then, the convex potentials are trained on this latent space with a small training sample size. In the generation phase, we randomly select an image z_k from each class $k \in \mathbb{K}$, encode it into the latent space, denoted by $\tilde{z}_k := \text{Encoder}(z_k)$, and then map $\tilde{z}_k$ by the gradient of the convex potential of the other class, i.e., $\nabla \tilde{\varphi}_{|1-k|}(\tilde{z}_k)$, and finally draw the attacked image $\text{Decoder}\left(\nabla \tilde{\varphi}_{|1-k|}(\tilde{z}_k)\right)$. Figure 6 illustrates the trajectory of injecting adversarial attacks on digits in both classes in 7 trials. As illustrated, our adversarial attack naturally achieves conversions of different classes of numbers through convex

gradient mapping in a few steps.

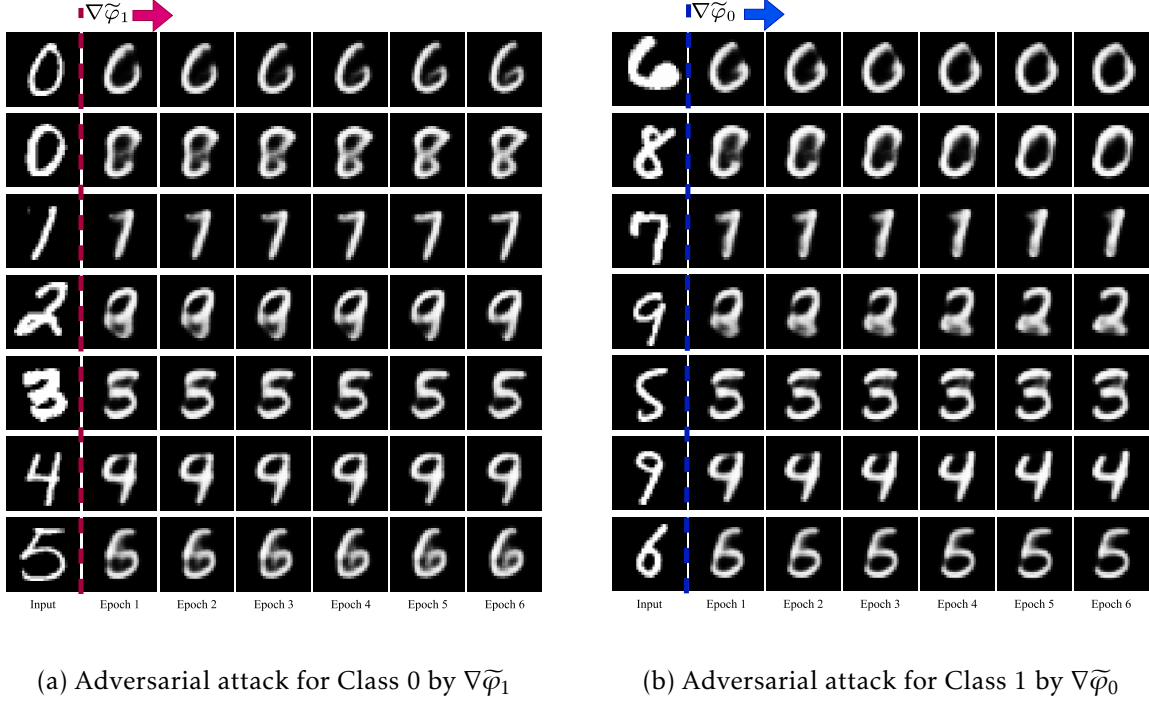


Figure 6: Adversarial attacks on digits by trained convex gradients. The first column displays the original images, while the following six columns show the attacked images from Epochs 1 to 6.

5.3 High-energy Physics Higgs Dataset

In this subsection, we evaluate our method on the Higgs dataset, which distinguishes between a signal process where new theoretical Higgs bosons (HIGGS) are produced, and a background process with the identical decay products but distinct kinematics [42]. We consider 21 original features and large training sample sizes.

Table 6 reports the accuracy and computation time of our method against other baselines over varying training sample sizes. As shown, the SDRO method achieves the highest accuracy in a relatively low computation time. Note that Table 6 does not include the result of WDRO, since when the sample size is large, WDRO did not produce an acceptable solution within the 1,800-second time limit. This is because WDRO requires solving a linear program with $\mathcal{O}(n_0 n_1)$ decision variables (see Appendix D) and becomes computationally prohibitive for large sample sizes.

In Table 7, we show that SDRO remains the highest accuracy over 70% across all non-zero attack intensities under both norms. This result demonstrates that our method can still maintain its accuracy even under significant adversarial perturbations compared to other baselines.

Table 6: Accuracy (%) and computation time (s) of methods on the Higgs dataset, averaged over observations $m \in [10]$ and 10 trials. Bold represents the highest accuracy in each row.

$n_0 = n_1$ ($\times 10^3$)	Accuracy (%)					Computation Time (s)				
	SDRO	FDRO	LR	SVM	3NN	SDRO	FDRO	LR	SVM	3NN
1	68.09	62.07	59.73	64.13	57.92	1.11	2.09	0.02	0.18	5.09
2	71.32	65.56	60.45	66.67	58.70	1.06	2.06	0.01	0.58	11.06
3	72.83	67.20	60.68	67.87	58.97	1.13	2.69	0.01	0.74	9.86
4	73.88	68.03	60.95	68.88	60.17	1.64	3.65	0.02	1.51	11.21
5	74.62	68.77	61.01	69.56	61.00	1.79	3.56	0.02	2.20	11.16
6	75.37	69.15	61.07	70.27	61.54	1.87	3.30	0.02	3.01	13.74
7	75.71	69.21	61.09	70.67	62.86	2.37	3.35	0.03	4.46	15.90
8	76.25	69.60	61.22	71.13	63.66	2.47	3.30	0.03	5.66	16.71
9	76.26	69.71	61.18	71.14	63.85	2.54	3.73	0.02	7.33	18.53
10	76.66	69.91	61.29	71.59	64.89	2.96	3.52	0.03	8.65	21.41

Note: The WDRO method did not produce an acceptable solution within the 1,800-second time limit.

Table 7: Robustness comparison on Higgs dataset after PGD attack, averaged over $m \in [10]$ and 5 trials (accuracy %, $n_0 = n_1 = 2,000$).

Method	PGD- L_2 Attack				PGD- L_∞ Attack			
	$\Delta = 0.025$	$\Delta = 0.05$	$\Delta = 0.075$	$\Delta = 0.1$	$\Delta = 0.025$	$\Delta = 0.05$	$\Delta = 0.075$	$\Delta = 0.1$
SDRO	72.342	71.998	71.887	71.891	73.352	72.758	72.436	72.284
FDRO	66.999	64.888	62.736	60.835	62.780	55.396	48.914	44.461
LR	62.945	60.913	59.193	57.582	59.708	54.046	48.528	44.412
SVM	67.669	65.849	64.191	62.716	64.524	58.350	53.250	49.614
3NN	58.204	54.669	50.782	48.015	50.264	38.658	30.418	26.266

5.4 Human Face Generation

In this subsection, similar to [28], we present a qualitative experiment on the human face dataset FFHQ [43] that illustrates how our method generates adversarial-attacked high-dimensional images.

We take the gender attribute of the FFHQ dataset as the binary label in the following image generation experiment. We labeled Class 0 as *Female* and Class 1 as *Male*. In the training phase, we first map the original 1024×1024 images to a 512-dimensional latent space by a pre-trained adversarial latent autoencoder [44]. Then, we train two HyCNNs by Algorithm 1 on the latent space. In the generation phase, we randomly select three images in each class. These samples are first encoded into a latent space, then mapped by our trained convex gradient, and finally recovered in the original space by the pre-trained decoder.

For all experiments, we set $n_0 = n_1 = 500$, $\lambda = 1$, $\epsilon = 0.1$, $T = 30$, and each HyCNN has a single

hidden layer with 128 units. Figure 7 shows the trajectories of generating the attacked images of all selected samples in the gender classification task. In each subfigure, the first column shows the original image, and the other 12 columns represent the images generated using the convex gradient maps trained for 30 epochs, respectively. As illustrated, the samples we generated underwent meaningful modifications to the original images, such that the generated images gradually acquired visual characteristics associated with the opposite dataset label after a small number of training epochs.

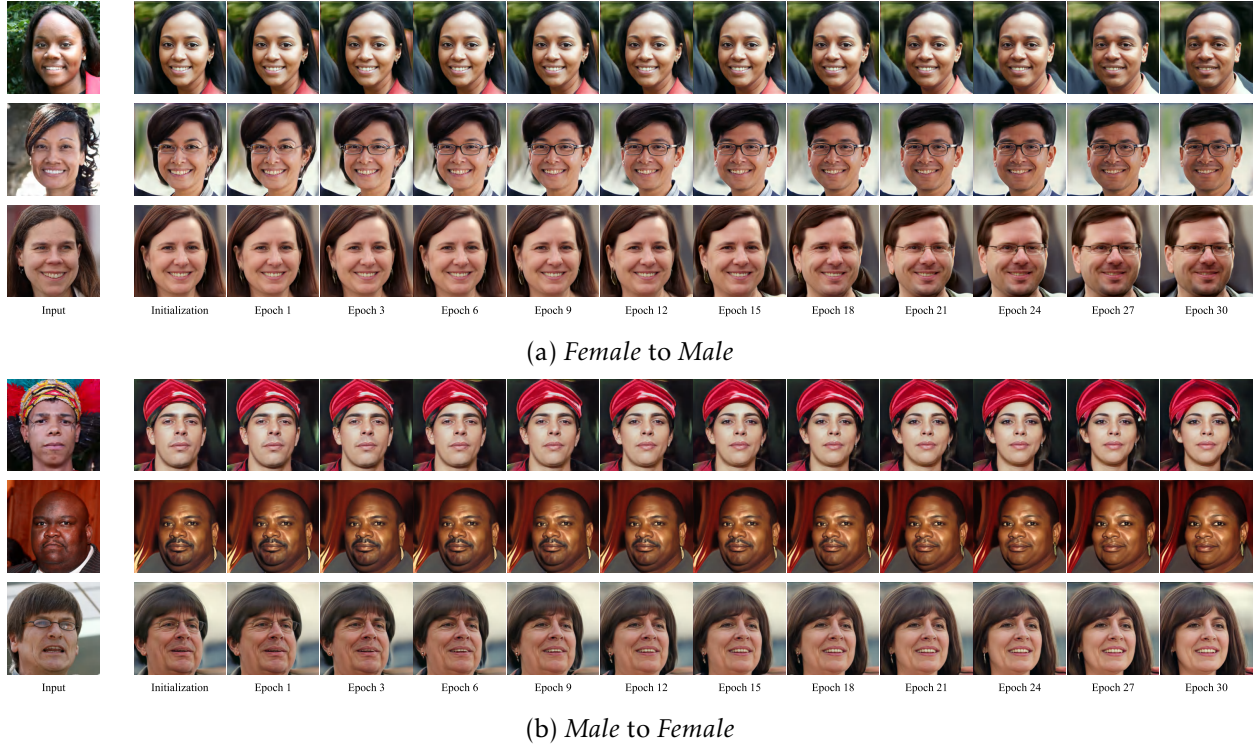


Figure 7: The trajectory of image generation between *Female* (Class 0) and *Male* (Class 1). The first column displays the original images, while the following columns show the images generated during 30 epochs.

6 Conclusions and Discussions

As a fundamental method in statistics, robust hypothesis testing addresses noisy data, distributional shifts, and measurement errors in real-world, uncertain environments. This paper develops a generative framework for Sinkhorn distributionally robust hypothesis testing (SDRHT), overcoming the non-continuity of the least-favorable distributions of WDRO-based testing and the computational challenges of existing SDRHT methods. By exploiting a conditional-KL representation of Sinkhorn ambiguity sets and strong duality, this framework reformulates the infinite-dimensional problem as an optimization over parametrized convex potentials. We further establish approximation guarantees for the potentials and their gradients, as well as the distributional

universality of the induced transport maps. Numerical results demonstrate the superior performance and robustness of our method compared to existing methods on tasks with different sample sizes and dimensions.

There are several topics worth investigating for future work. First, it is interesting to extend our method to multi-class hypothesis testing, exploring its closed-form optimal detector and efficient algorithms. Then, since the transport map between distributions is the gradient of a convex potential, directly parameterizing the gradient instead of the potential may yield a better approximation of the transport map. Finally, for distributions with extremely high dimensions, e.g., high-resolution images, it is important to develop methods that efficiently solve transport maps while maintaining the convexity of potentials.

References

- [1] J. Neyman and E. S. Pearson, “On the problem of the most efficient tests of statistical hypotheses,” *Philosophical Transactions of the Royal Society of London. Series A, Containing Papers of a Mathematical or Physical Character*, vol. 231, no. 694-706, pp. 289–337, 1933.
- [2] B. C. Levy, *Principles of signal detection and parameter estimation*. Springer Science & Business Media, 2008.
- [3] P. Schober and T. R. Vetter, “Two-sample unpaired t-tests in medical research,” *Anesthesia & Analgesia*, vol. 129, no. 4, p. 911, 2019.
- [4] V. Chandola, A. Banerjee, and V. Kumar, “Anomaly detection: A survey,” *ACM Computing Surveys (CSUR)*, vol. 41, no. 3, pp. 1–58, 2009.
- [5] P. J. Huber, “A robust version of the probability ratio test,” *The Annals of Mathematical Statistics*, pp. 1753–1758, 1965.
- [6] G. Gül and A. M. Zoubir, “Minimax robust hypothesis testing,” *IEEE Transactions on Information Theory*, vol. 63, no. 9, pp. 5572–5587, 2017.
- [7] D. Kuhn, S. Shafiee, and W. Wiesemann, “Distributionally robust optimization,” *Acta Numerica*, vol. 34, pp. 579–804, 2025.
- [8] R. Gao, L. Xie, Y. Xie, and H. Xu, “Robust hypothesis testing using Wasserstein uncertainty sets,” *Advances in Neural Information Processing Systems*, vol. 31, 2018.
- [9] P. Mohajerin Esfahani and D. Kuhn, “Data-driven distributionally robust optimization using the Wasserstein metric: Performance guarantees and tractable reformulations,” *Mathematical Programming*, vol. 171, no. 1, pp. 115–166, 2018.
- [10] P. J. Huber, *Robust statistical procedures*. SIAM, 1996.
- [11] B. C. Levy, “Robust hypothesis testing with a relative entropy tolerance,” *IEEE Transactions on Information Theory*, vol. 55, no. 1, pp. 413–421, 2008.

- [12] 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*, vol. 59, no. 2, pp. 341–357, 2013.
- [13] E. Delage and Y. Ye, “Distributionally robust optimization under moment uncertainty with application to data-driven problems,” *Operations Research*, vol. 58, no. 3, pp. 595–612, 2010.
- [14] W. Wiesemann, D. Kuhn, and M. Sim, “Distributionally robust convex optimization,” *Operations Research*, vol. 62, no. 6, pp. 1358–1376, 2014.
- [15] Z. Sun and S. Zou, “A data-driven approach to robust hypothesis testing using kernel MMD uncertainty sets,” in *2021 IEEE International Symposium on Information Theory (ISIT)*. IEEE, 2021, pp. 3056–3061.
- [16] J.-J. Zhu, W. Jitkrittum, M. Diehl, and B. Schölkopf, “Kernel distributionally robust optimization: Generalized duality theorem and stochastic approximation,” in *International Conference on Artificial Intelligence and Statistics*. PMLR, 2021, pp. 280–288.
- [17] L. Xie, R. Gao, and Y. Xie, “Robust hypothesis testing with Wasserstein uncertainty sets,” *arXiv preprint arXiv:2105.14348*, 2021.
- [18] J. Wang, R. Gao, and Y. Xie, “Sinkhorn distributionally robust optimization,” *Operations Research*, 2025.
- [19] F. Zhang and J. Wang, “Contextual distributionally robust optimization with causal and continuous structure: An interpretable and tractable approach,” *arXiv preprint arXiv:2601.11016*, 2026.
- [20] G. Peyré and M. Cuturi, *Computational optimal transport: With applications to data science*. Now Foundations and Trends, 2019.
- [21] J. Wang and Y. Xie, “A data-driven approach to robust hypothesis testing using Sinkhorn uncertainty sets,” in *2022 IEEE International Symposium on Information Theory (ISIT)*. IEEE, 2022, pp. 3315–3320.
- [22] J. Wang, R. Gao, and Y. Xie, “Non-convex robust hypothesis testing using Sinkhorn uncertainty sets,” in *2024 IEEE International Symposium on Information Theory (ISIT)*. IEEE, 2024, pp. 843–848.
- [23] C.-H. Lai, Y. Song, D. Kim, Y. Mitsufuji, and S. Ermon, “The principles of diffusion models,” *arXiv preprint arXiv:2510.21890*, 2025.
- [24] Y. Lipman, R. T. Chen, H. Ben-Hamu, M. Nickel, and M. Le, “Flow matching for generative modeling,” *arXiv preprint arXiv:2210.02747*, 2022.
- [25] T. Dai, D. Simchi-Levi, M. X. Wu, and Y. Xie, “Assured autonomy: How operations research powers and orchestrates generative AI systems,” *arXiv preprint arXiv:2512.23978*, 2025.

- [26] C. Xu, J. Lee, X. Cheng, and Y. Xie, “Flow-based distributionally robust optimization,” *IEEE Journal on Selected Areas in Information Theory*, 2024.
- [27] J. Wang, “Iterative sampling methods for sinkhorn distributionally robust optimization,” *arXiv preprint arXiv:2512.12550*, 2025.
- [28] Z. Xu and J.-J. Zhu, “Gradient flow sampler-based distributionally robust optimization,” *arXiv preprint arXiv:2510.25956*, 2025.
- [29] L. Zhu and Y. Xie, “Distributionally robust optimization via iterative algorithms in continuous probability spaces,” *arXiv preprint arXiv:2412.20556*, 2024.
- [30] X. Cheng, Y. Xie, L. Zhu, and Y. Zhu, “Worst-case generation via minimax optimization in Wasserstein space,” *arXiv preprint arXiv:2512.08176*, 2025.
- [31] A. Nemirovski and A. Shapiro, “Convex approximations of chance-constrained programs,” *SIAM Journal on Optimization*, vol. 17, no. 4, pp. 969–996, 2007.
- [32] S. Chewi, J. Niles-Weed, and P. Rigollet, *Statistical optimal transport*. Springer, 2025.
- [33] S. Ubaru, J. Chen, and Y. Saad, “Fast estimation of $\text{tr}(f(a))$ via stochastic lanczos quadrature,” *SIAM Journal on Matrix Analysis and Applications*, vol. 38, no. 4, pp. 1075–1099, 2017.
- [34] C.-W. Huang, R. T. Chen, C. Tsirigotis, and A. Courville, “Convex potential flows: Universal probability distributions with optimal transport and convex optimization,” *arXiv preprint arXiv:2012.05942*, 2020.
- [35] B. Amos, L. Xu, and J. Z. Kolter, “Input convex neural networks,” in *International Conference on Machine Learning*. PMLR, 2017, pp. 146–155.
- [36] S. Hundrieser, I. Kong, and J. Schmidt-Hieber, “Hyper input convex neural networks for shape constrained learning and optimal transport,” *arXiv preprint arXiv:2604.26942*, 2026.
- [37] I. Goodfellow, D. Warde-Farley, M. Mirza, A. Courville, and Y. Bengio, “Maxout networks,” in *International Conference on Machine Learning*. PMLR, 2013, pp. 1319–1327.
- [38] M. F. Hutchinson, “A stochastic estimator of the trace of the influence matrix for laplacian smoothing splines,” *Communications in Statistics-Simulation and Computation*, vol. 18, no. 3, pp. 1059–1076, 1989.
- [39] A.-A. Pooladian and J. Niles-Weed, “Entropic estimation of optimal transport maps,” *arXiv preprint arXiv:2109.12004*, 2021.
- [40] Y. Chen, Y. Shi, and B. Zhang, “Optimal control via neural networks: A convex approach,” *arXiv preprint arXiv:1805.11835*, 2018.
- [41] A. Madry, A. Makelov, L. Schmidt, D. Tsipras, and A. Vladu, “Towards deep learning models resistant to adversarial attacks,” *arXiv preprint arXiv:1706.06083*, 2017.

- [42] P. Baldi, P. Sadowski, and D. Whiteson, “Searching for exotic particles in high-energy physics with deep learning,” *Nature communications*, vol. 5, no. 1, p. 4308, 2014.
- [43] T. Karras, S. Laine, and T. Aila, “A style-based generator architecture for generative adversarial networks,” in *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 2019, pp. 4401–4410.
- [44] S. Pidhorskyi, D. A. Adjeroh, and G. Doretto, “Adversarial latent autoencoders,” in *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 2020, pp. 14 104–14 113.
- [45] R. Cescon, A. Martin, and G. Ferrari-Trecate, “Sinkhorn ambiguity sets for distributionally robust control: Convexity, weak compactness, and tractability,” *arXiv preprint arXiv:2605.03845*, 2026.
- [46] S. Shafiee, L. Aolaritei, F. Dörfler, and D. Kuhn, “Nash equilibria, regularization, and computation in optimal transport-based distributionally robust optimization,” *Operations Research*, 2025.
- [47] M. D. Donsker and S. S. Varadhan, “Asymptotic evaluation of certain markov process expectations for large time. iv,” *Communications on pure and applied mathematics*, vol. 36, no. 2, pp. 183–212, 1983.
- [48] A. Shapiro, D. Dentcheva, and A. Ruszczyński, *Lectures on stochastic programming: modeling and theory*. SIAM, 2021.

Appendix

A Summary of Notations

For integer $K \in \mathbb{Z}_+$, define $[K] := \{1, \dots, K\}$, which is the set of positive running indices to K . The list of important sets, spaces, and notations used in this paper is shown in Table 8. The list of important functions is shown in Table 9.

Table 8: Notations List

Spaces	Description
Ω	Closed sample space, $\Omega \subseteq \mathbb{R}^d$.
$\mathcal{F}$	Functional space, which is not necessarily compact.
Basic Sets	
$\mathbb{K}$	Indices set of hypotheses, $\mathbb{K} = \{0, 1\}$.
$\mathcal{P}(\Omega)$	The set of all probability distributions on Ω .
$\mathcal{M}(\Omega)$	The set of measures on Ω .
$\Gamma(\mathbb{P}, \mathbb{Q})$	Joint distribution set of marginal distribution $\mathbb{P}$ and $\mathbb{Q}$.
$\mathcal{P}_k$	The distribution set under hypotheses H_k , $\mathcal{P}_k \subset \mathcal{P}(\Omega)$ for each $k \in \mathbb{K}$.
$\mathcal{F}_{\text{CVX}}$	Set of convex functions $\varphi : \mathbb{R}^d \rightarrow \mathbb{R} \cup \{+\infty\}$.
Parameters	
ϵ	Entropic regularization parameter in Sinkhorn discrepancy, $\epsilon \geq 0$.
ρ_k	Radius for Sinkhorn-based ambiguity set, for each $k \in \mathbb{K}$.
λ_k	Penalty parameter in soft-constraint problem (soft-SDRHT), $\lambda_k \geq 0$ for each $k \in \mathbb{K}$.
Distributions	
$\widehat{\mathbb{P}}_k$	Empirical distribution with n_k samples for each $k \in \mathbb{K}$.
$\mathcal{K}_{\mathbf{x}_k^{(i)}, \epsilon}$	Kernel extension for any sample point $i \in [n_k]$ and each $k \in \mathbb{K}$.
$\mathbb{P}_k^\epsilon$	Mixture of continuous distributions $\mathcal{K}_{\mathbf{x}_k^{(i)}, \epsilon}$ for all $i \in [n_k]$, see Eq. (4).
$\widetilde{\mathbb{P}}_k$	Discrete approximation of distribution $\mathbb{P}_k^\epsilon$ with N samples, for each $k \in \mathbb{K}$.
Discrepancies	
$\mathcal{S}_\epsilon(\mathbb{P}, \mathbb{Q})$	Sinkhorn discrepancy from distribution $\mathbb{P}$ to $\mathbb{Q}$, see Definition 2.
$\mathcal{D}_{\text{KL}}(\mathbb{P}, \mathbb{Q})$	Kullback–Leibler divergence (KL divergence) from distribution $\mathbb{P}$ to $\mathbb{Q}$.
Ambiguity Sets	
$\mathbb{B}_{\rho_k}(\widehat{\mathbb{P}}_k)$	Sinkhorn discrepancy-based ambiguity set centered at $\widehat{\mathbb{P}}_k$ for each $k \in \mathbb{K}$, see Eq. (1).
$\mathbb{B}_{\widetilde{\rho}_k}(\mathbb{P}_k^\epsilon)$	Equivalent KL-based ambiguity set of $\mathbb{B}_{\rho_k}(\widehat{\mathbb{P}}_k)$ centered at $\mathbb{P}_k^\epsilon$, see Proposition 1.

Table 9: Functions List

Functions	Description
$c : \Omega \times \Omega \rightarrow \mathbb{R}$	Transport cost function satisfying Assumption 1, we choose $c(\mathbf{x}, \mathbf{z}) := \frac{1}{2} \ \mathbf{x} - \mathbf{z}\ _2^2$.
$\ell : \mathbb{R} \rightarrow \mathbb{R}_+ \cup \{\infty\}$	Generating function (see Definition 1), we choose $\ell(t) = \exp(t)$.
$\phi : \Omega \rightarrow \mathbb{R}$	Detector for hypothesis testing, which is measurable and upper semicontinuous.
$\phi_{\text{opt}} : \Omega \rightarrow \mathbb{R}$	Optimal detector for generating function $\ell(t) = \exp(t)$, $\phi_{\text{opt}} = \frac{1}{2} \log(p_0/p_1)$.
$\varphi_1, \varphi_2 : \mathbb{R}^d \rightarrow \mathbb{R}$	Convex functions, decision variables in the problem (7).
$\varphi^* : \mathbb{R}^d \rightarrow \mathbb{R} \cup \{\infty\}$	Convex conjugate of a convex function $\varphi : \mathbb{R}^d \rightarrow \mathbb{R} \cup \{\infty\}$.
$\text{HyCNN}_k : \mathbb{R}^d \rightarrow \mathbb{R}$	Convex function formed by the HyCNN for each $k \in \mathbb{K}$.
$\widetilde{\varphi}_k : \mathbb{R}^d \rightarrow \mathbb{R}$	Strongly convex function, $\widetilde{\varphi}_k(\mathbf{z}) = \text{HyCNN}_k(\mathbf{z}) + \frac{\alpha}{2} \ \mathbf{z}\ _2^2$ for each $k \in \mathbb{K}$ and $\alpha > 0$.
$L_k : \Omega \rightarrow \mathbb{R}$	Loss function of a given detector ϕ for sample ω , satisfying Assumption 2.
$\Phi : \mathcal{P}(\Omega)^2 \rightarrow \mathbb{R}$	Convex surrogate risk of the detector ϕ .
$\underline{g}_{\epsilon, k}^c$	Continuous extension for the optimal entropic potential $g_{\epsilon, k, (n_k, N)}$ in Algorithm 4.
$\widehat{\mathcal{S}}_{\epsilon, (n_k, N)}$	Finite sample estimator for the Sinkhorn discrepancy.

B Technical Analyses and Proofs

B.1 Proof of Proposition 1

Proof of Proposition 1. By the definition of Sinkhorn discrepancy, for distributions $\mathbb{P}, \mathbb{Q} \in \mathcal{P}(\Omega)$, it is equivalent to

$$\mathcal{S}_\epsilon(\mathbb{P}, \mathbb{Q}) = \inf_{\gamma \in \Gamma(\mathbb{P}, \mathbb{Q})} \int_{\Omega \times \Omega} \left(c(\mathbf{x}, \mathbf{z}) + \epsilon \cdot \log \frac{d\gamma(\mathbf{x}, \mathbf{z})}{d\mathbb{P}(\mathbf{x})d\mathbb{Q}(\mathbf{z})} \right) d\gamma(\mathbf{x}, \mathbf{z}). \quad (20)$$

Taking the reference distribution as the empirical distribution $\widehat{\mathbb{P}} \in \mathcal{P}(\Omega)$, Eq. (20) can be further reformulated to

$$\begin{aligned} \mathcal{S}_\epsilon(\widehat{\mathbb{P}}, \mathbb{Q}) &= \inf_{\gamma \in \Gamma(\widehat{\mathbb{P}}, \mathbb{Q})} \mathbb{E}_{\mathbf{x} \sim \widehat{\mathbb{P}}} \left[\int_{\Omega} \left(c(\mathbf{x}, \mathbf{z}) + \epsilon \log \frac{d\gamma_{\mathbf{x}}(\mathbf{z})}{d\mathbb{Q}(\mathbf{z})} \right) d\mathbb{Q}(\mathbf{z}) \right] \\ &= \inf_{\gamma \in \Gamma(\widehat{\mathbb{P}}, \mathbb{Q})} \epsilon \cdot \mathbb{E}_{\mathbf{x} \sim \widehat{\mathbb{P}}} \left[\int_{\Omega} \log \frac{d\gamma_{\mathbf{x}}(\mathbf{z})}{e^{-c(\mathbf{x}, \mathbf{z})} \cdot d\mathbb{Q}(\mathbf{z})} d\mathbb{Q}(\mathbf{z}) \right] \\ &= \inf_{\gamma \in \Gamma(\widehat{\mathbb{P}}, \mathbb{Q})} \epsilon \cdot \mathbb{E}_{\mathbf{x} \sim \widehat{\mathbb{P}}} \left[\int_{\Omega} \log \frac{d\gamma_{\mathbf{x}}(\mathbf{z})}{d\mathcal{K}_{\mathbf{x}, \epsilon}(\mathbf{z})} d\mathbb{Q}(\mathbf{z}) - \log \int_{\mathbb{R}^d} e^{-c(\mathbf{x}, \mathbf{z})/\epsilon} d\mathbb{Q}(\mathbf{z}) \right] \\ &= \inf_{\gamma \in \Gamma(\widehat{\mathbb{P}}, \mathbb{Q})} \epsilon \cdot \mathbb{E}_{\mathbf{x} \sim \widehat{\mathbb{P}}} \left[\mathcal{D}_{\text{KL}}(\gamma_{\mathbf{x}} \parallel \mathcal{K}_{\mathbf{x}, \epsilon}) \right] - \epsilon \cdot \mathbb{E}_{\mathbf{x} \sim \widehat{\mathbb{P}}} \left[\log \int_{\mathbb{R}^d} e^{-c(\mathbf{x}, \mathbf{z})/\epsilon} d\mathbb{Q}(\mathbf{z}) \right], \end{aligned} \quad (21)$$

where

$$d\mathcal{K}_{\mathbf{x}, \epsilon}(\mathbf{z}) := \frac{e^{-c(\mathbf{x}, \mathbf{z})/\epsilon}}{\int_{\mathbb{R}^d} e^{-c(\mathbf{x}, \mathbf{u})/\epsilon} \cdot d\mathbb{Q}(\mathbf{u})} \cdot d\mathbb{Q}(\mathbf{z}).$$

Thus, the ambiguity sets in Eq. (1) is equivalent to

$$\inf_{\gamma \in \Gamma(\widehat{\mathbb{P}}_k, \mathbb{Q})} \epsilon \cdot \mathbb{E}_{\mathbf{x} \sim \widehat{\mathbb{P}}_k} \left[\mathcal{D}_{\text{KL}}(\gamma_{\mathbf{x}} \parallel \mathcal{K}_{\mathbf{x}, \epsilon}) \right] \leq \rho_k + \epsilon \cdot \mathbb{E}_{\mathbf{x} \sim \widehat{\mathbb{P}}_k} \left[\log \int_{\mathbb{R}^d} e^{-c(\mathbf{x}, \mathbf{z})/\epsilon} d\mathbb{Q}(\mathbf{z}) \right], \quad \forall k \in \mathbb{K},$$

that is,

$$\mathbb{B}_{\bar{\rho}_k}(\mathbb{P}_k^\epsilon) := \{\mathbb{Q} \in \mathcal{P}(\Omega) : \inf_{\gamma \in \Gamma(\widehat{\mathbb{P}}_k, \mathbb{Q})} \mathbb{E}_{\mathbf{x} \sim \widehat{\mathbb{P}}_k} [\mathcal{D}_{\text{KL}}(\gamma_{\mathbf{x}} \| \mathcal{K}_{\mathbf{x}, \epsilon})] \leq \bar{\rho}_k / \epsilon\}, \quad \forall k \in \mathbb{K},$$

where $\bar{\rho}_k := \rho_k + \epsilon \mathbb{E}_{\mathbf{x} \sim \widehat{\mathbb{P}}_k} \left[\log \int_{\mathbb{R}^d} e^{-c(\mathbf{x}, \mathbf{z})/\epsilon} d\mathbf{z} \right]$ and probability density $d\mathbb{P}_k^\epsilon(\mathbf{z}) := \mathbb{E}_{\mathbf{x} \sim \widehat{\mathbb{P}}_k} [d\mathcal{K}_{\mathbf{x}, \epsilon}(\mathbf{z})]$. This completes the proof. $\square$

B.2 Proof of Theorem 1

To prove Theorem 1, we first establish the weak compactness of the Sinkhorn ambiguity sets and the weak upper semicontinuity of the expected loss.

Based on Eq. (3), the following lemma holds.

Lemma 1 (Weak Compactness of the Ambiguity Sets). *Under Assumption 1, the ambiguity set defined in (3) is weakly compact in $\mathcal{P}(\Omega)$ for any closed set $\Omega \subseteq \mathbb{R}^d$, $k \in \mathbb{K}$, and $\bar{\rho}_k \geq 0$.*

The proof of Lemma 1 is provided in the following Appendix B.2.1. While [45] also confirms the weak compactness of Sinkhorn ambiguity sets based on the properties of Wasserstein balls, our proof follows directly from the KL representation in Proposition 1 and holds under more general assumptions.

Lemma 2 (Weak Upper Semicontinuity of the Expected Loss). *Under Assumption 2(I), $\mathbb{E}_{\omega \sim \mathbb{Q}}[L_k(\phi, \omega)]$ is weakly upper semicontinuous in $\mathbb{Q} \in \mathbb{B}_{\bar{\rho}_k}(\mathbb{P}_k^\epsilon)$ for any $k \in \mathbb{K}$ and $\phi \in \mathcal{F}$.*

The proof of Lemma 2 is provided in the following Appendix B.2.2.

B.2.1 Proof of Lemma 1

Proof of Lemma 1. To prove $\mathbb{B}_{\bar{\rho}_k}(\mathbb{P}_k^\epsilon)$ is weakly compact, we first show its connection with KL-divergence-based ambiguity set. Since the KL-divergence is a convex function, by Jensen's inequality, we have

$$\begin{aligned} \mathcal{D}_{\text{KL}}(\mathbb{Q} \| \mathbb{P}_k^\epsilon) &\leq \inf_{\gamma \in \Gamma(\widehat{\mathbb{P}}_k, \mathbb{Q})} \mathbb{E}_{\mathbf{x} \sim \widehat{\mathbb{P}}_k} [\mathcal{D}_{\text{KL}}(\gamma_{\mathbf{x}} \| \mathcal{K}_{\mathbf{x}, \epsilon})] \\ &\leq \frac{\rho_k}{\epsilon} + \mathbb{E}_{\mathbf{x} \sim \widehat{\mathbb{P}}_k} \left[\log \int_{\mathbb{R}^d} e^{-c(\mathbf{x}, \mathbf{z})/\epsilon} d\mathbf{z} \right], \quad \forall k \in \mathbb{K}, \end{aligned}$$

where probability distribution $\mathbb{Q} \in \mathcal{P}(\Omega)$. Let

$$\mathbb{B}'_{\bar{\rho}_k}(\mathbb{P}_k^\epsilon) := \{\mathbb{Q} : \mathcal{D}_{\text{KL}}(\mathbb{Q} \| \mathbb{P}_k^\epsilon) \leq \bar{\rho}_k / \epsilon\}, \quad \forall k \in \mathbb{K}.$$

According to Proposition 3.12 in [7], the KL-divergence-based ambiguity sets $\mathbb{B}'_{\bar{\rho}_k}(\mathbb{P}_k^\epsilon)$ for any $k \in \mathbb{K}$ are weakly compact, for any closed set $\Omega \subseteq \mathbb{R}^d$ and $k \in \mathbb{K}$, distribution $\mathbb{P}_k^\epsilon \in \mathcal{P}(\Omega)$ and $\bar{\rho}_k \geq 0$.

Since $\mathbb{B}_{\bar{\rho}_k}(\mathbb{P}_k^\epsilon) \subseteq \mathbb{B}'_{\bar{\rho}_k}(\mathbb{P}_k^\epsilon)$ for any $k \in \mathbb{K}$, it suffices to prove that $\mathbb{B}_{\bar{\rho}_k}(\mathbb{P}_k^\epsilon)$ is weakly closed, as weakly closed subsets of weakly compact sets are weakly compact [46, Lemma 1]. Define

$$dM_{k,\epsilon}(\mathbf{x}, \mathbf{z}) := d\widehat{\mathbb{P}}_k(\mathbf{x}) \cdot d\mathcal{K}_{\mathbf{x},\epsilon}(\mathbf{z}) = \frac{e^{-c(\mathbf{x},\mathbf{z})/\epsilon}}{\int_{\mathbb{R}^d} e^{-c(\mathbf{x},\mathbf{u})/\epsilon} \cdot d\mathbf{u}} \cdot d\widehat{\mathbb{P}}_k(\mathbf{x}) \cdot dv(\mathbf{z}), \quad \forall k \in \mathbb{K},$$

According to Eq. (20), we have

$$\inf_{\gamma \in \Gamma(\widehat{\mathbb{P}}_k, \mathbb{Q})} \mathbb{E}_{\mathbf{x} \sim \widehat{\mathbb{P}}_k} \left[\mathcal{D}_{\text{KL}}(\gamma_{\mathbf{x}} \| \mathcal{K}_{\mathbf{x},\epsilon}) \right] = \inf_{\gamma \in \Gamma(\widehat{\mathbb{P}}_k, \mathbb{Q})} \mathcal{D}_{\text{KL}}(\gamma \| M_{k,\epsilon}), \quad \forall k \in \mathbb{K}.$$

According to Proposition 3.12 in [7], to prove the Sinkhorn-based ambiguity set is weakly closed, it suffices to prove that function $f_k(\mathbb{Q}) := \inf_{\gamma \in \Gamma(\widehat{\mathbb{P}}_k, \mathbb{Q})} \mathcal{D}_{\text{KL}}(\gamma \| M_{k,\epsilon})$ is weakly lower semicontinuous in $\mathbb{Q}$ for each $k \in \mathbb{K}$, which implies that any sublevel set of $f(\mathbb{Q})$, i.e., $\{\mathbb{Q} : f(\mathbb{Q}) \leq c\}$ for any $c \in \mathbb{R}$, is weakly closed. Let $\{\mathbb{Q}_j\}_{j \in \mathbb{N}} \subset \mathcal{P}(\Omega)$ be any sequence of distributions that converges weakly to $\mathbb{Q}$. By the definition of infimum, there exists a sequence of couplings $\gamma_n \in \Gamma(\widehat{\mathbb{P}}_k, \mathbb{Q}_n)$ that converges weakly to $\gamma \in \Gamma(\widehat{\mathbb{P}}_k, \mathbb{Q})$ such that

$$\mathcal{D}_{\text{KL}}(\gamma_n \| M_{k,\epsilon}) \leq f_k(\mathbb{Q}_n) + \frac{1}{n}, \quad \forall k \in \mathbb{K}. \quad (22)$$

Then, for each $k \in \mathbb{K}$, we have

$$f_k(\mathbb{Q}) \leq \mathcal{D}_{\text{KL}}(\gamma \| M_{k,\epsilon}) \leq \liminf_{j \in \mathbb{N}} \mathcal{D}_{\text{KL}}(\gamma_j \| M_{k,\epsilon}) \leq \liminf_{j \in \mathbb{N}} f_k(\mathbb{Q}_j), \quad \forall k \in \mathbb{K},$$

where the second inequality is due to the weak lower semi-continuity of $\mathcal{D}_{\text{KL}}(\gamma \| M_{k,\epsilon})$ in γ [7, Proposition 3.12], and the final inequality is due to Eq. (22). Therefore, for each $k \in \mathbb{K}$, the function $f_k(\mathbb{Q})$ is weakly lower semicontinuous in $\mathbb{Q}$ which implies that the ambiguity set $\mathbb{B}_{\bar{\rho}_k}(\mathbb{P}_k^\epsilon)$ is weakly closed, and thereby weakly compact. This completes the proof. $\square$

B.2.2 Proof of Lemma 2

Proof of Lemma 2. We prove Lemma 2 in two steps. Specifically, we prove that the loss function L_k is uniformly integrable and $\mathbb{E}_{\omega \sim \mathbb{Q}}[L_k(\phi, \omega)]$ is weakly upper semicontinuous with respect to the ambiguity set $\mathbb{Q} \in \mathbb{B}_{\bar{\rho}_k}(\mathbb{P}_k^\epsilon)$ for each $k \in \mathbb{K}$ and any fixed ϕ , respectively.

Step 1 of Lemma 2: Based on the Donsker-Varadhan variational representation of the KL divergence [47], for any measurable function $f : \Omega \rightarrow \mathbb{R}$ and probability distribution $\mathbb{Q}$, we have

$$\mathbb{E}_{\omega \sim \mathbb{Q}}[f] - \log \mathbb{E}_{\omega \sim \mathbb{P}_k^\epsilon}[e^f] \leq \mathcal{D}_{\text{KL}}(\mathbb{Q}, \mathbb{P}_k^\epsilon) \leq \inf_{\gamma \in \Gamma(\widehat{\mathbb{P}}_k, \mathbb{Q})} \mathbb{E}_{\mathbf{x} \sim \widehat{\mathbb{P}}_k} \left[\mathcal{D}_{\text{KL}}(\gamma_{\mathbf{x}} \| \mathcal{K}_{\mathbf{x},\epsilon}) \right] \leq \frac{\bar{\rho}_k}{\epsilon}, \quad \forall k \in \mathbb{K}. \quad (23)$$

According to the Donsker-Varadhan variational representation part in relation (23), under Assumption 2(I), for $\beta_k > 0$, we have

$$\mathbb{E}_{\omega \sim \mathbb{Q}}[L_k(\phi, \omega)] \leq \frac{1}{\beta} \left(\mathcal{D}_{\text{KL}}(\mathbb{Q}, \mathbb{P}_k^\epsilon) + \log \mathbb{E}_{\omega \sim \mathbb{P}_k^\epsilon}[e^{\beta_k \cdot L_k(\phi, \omega)}] \right) < \infty, \quad \forall k \in \mathbb{K}, \quad (24)$$

which implies that the inner problem of (SDRHT) is bounded. For a truncation threshold $M > 0$, let $\mathbb{1}_{\{L_k(\phi, \omega) > M\}}$ be an indicator function that equals to 1 if $L_k(\phi, \omega) > M$ and 0 otherwise. Based on relation (23), we also have

$$\mathbb{E}_{\omega \sim \mathbb{Q}}[L_k(\phi, \omega) \cdot \mathbb{1}_{\{L_k(\phi, \omega) > M\}}] \leq \frac{1}{\beta} \left(\mathcal{D}_{\text{KL}}(\mathbb{Q}, \mathbb{P}_k^\epsilon) + \log \mathbb{E}_{\omega \sim \mathbb{P}_k^\epsilon} \left[e^{\beta_k \cdot L_k(\phi, \omega) \cdot \mathbb{1}_{\{L_k(\phi, \omega) > M\}}} \right] \right), \quad \forall k \in \mathbb{K},$$

and it follows that

$$\sup_{\mathbb{Q} \in \mathbb{B}_{\bar{\rho}_k}(\mathbb{P}_k^\epsilon)} \mathbb{E}_{\omega \sim \mathbb{Q}}[L_k(\phi, \omega) \cdot \mathbb{1}_{\{L_k(\phi, \omega) > M\}}] \leq \frac{\bar{\rho}_k}{\epsilon \beta} + \frac{1}{\beta} \log \mathbb{E}_{\omega \sim \mathbb{P}_k^\epsilon} \left[e^{\beta_k \cdot L_k(\phi, \omega) \cdot \mathbb{1}_{\{L_k(\phi, \omega) > M\}}} \right], \quad \forall k \in \mathbb{K}.$$

When $M \rightarrow \infty$, we have $\mathbb{1}_{\{L_k(\phi, \omega) > M\}} \rightarrow 0$ for $\mathbb{P}_k^\epsilon$ -almost every ω . Moreover, as both β_k and $L_k(\phi, \omega)$ are non-negative, we have

$$e^{\beta_k \cdot L_k(\phi, \omega) \cdot \mathbb{1}_{\{L_k(\phi, \omega) > M\}}} \leq e^{\beta_k \cdot L_k(\phi, \omega)}, \quad \forall k \in \mathbb{K},$$

where $e^{\beta_k \cdot L_k(\phi, \omega)}$ is integrable under Assumption 2(I). Hence, by the dominated convergence theorem,

$$\lim_{M \rightarrow \infty} \log \mathbb{E}_{\omega \sim \mathbb{P}_k^\epsilon} \left[e^{\beta_k \cdot L_k(\phi, \omega) \cdot \mathbb{1}_{\{L_k(\phi, \omega) > M\}}} \right] = \log \mathbb{E}_{\omega \sim \mathbb{P}_k^\epsilon} [e^0] = 0, \quad \forall k \in \mathbb{K}.$$

Thus, we have

$$\lim_{M \rightarrow \infty} \sup_{\mathbb{Q} \in \mathbb{B}_{\bar{\rho}_k}(\mathbb{P}_k^\epsilon)} \mathbb{E}_{\omega \sim \mathbb{Q}}[L_k(\phi, \omega) \cdot \mathbb{1}_{\{L_k(\phi, \omega) > M\}}] \leq \frac{\bar{\rho}_k}{\epsilon \beta_k}, \quad \forall k \in \mathbb{K}. \quad (25)$$

Since Assumption 2(I) guarantees that this property holds for any $\beta_k > 0$, let $\beta_k \rightarrow \infty$, and then we have

$$\lim_{M \rightarrow \infty} \sup_{\mathbb{Q} \in \mathbb{B}_{\bar{\rho}_k}(\mathbb{P}_k^\epsilon)} \mathbb{E}_{\omega \sim \mathbb{Q}}[L_k(\phi, \omega) \cdot \mathbb{1}_{\{L_k(\phi, \omega) > M\}}] \leq \inf_{\beta > 0} \frac{\bar{\rho}_k}{\epsilon \beta} = 0, \quad \forall k \in \mathbb{K}, \quad (26)$$

which completes the proof of uniform integrability.

Step 2 of Lemma 2: We next prove the weak upper semicontinuity of $\mathbb{E}_{\omega \sim \mathbb{Q}}[L_k(\phi, \omega)]$ in $\mathbb{Q} \in \mathbb{B}_{\bar{\rho}_k}(\mathbb{P}_k^\epsilon)$. We consider a continuous detector ϕ and a non-negative, non-decreasing, convex, and continuous generating function ℓ . Then, the loss functions L_0 and L_1 are non-negative and upper semicontinuous for all $k \in \mathbb{K}$. By Lemma 1, the ambiguity set $\mathbb{B}_{\bar{\rho}_k}(\mathbb{P}_k^\epsilon)$ is weakly closed. For each $k \in \mathbb{K}$, let $\mathbb{Q}_j \in \mathbb{B}_{\bar{\rho}_k}(\mathbb{P}_k^\epsilon)$ for all $j \in \mathbb{N}$ be any sequence of distributions that converges weakly to $\mathbb{Q}$, and then we have

$$\begin{aligned} & \limsup_{j \in \mathbb{N}} \mathbb{E}_{\omega \sim \mathbb{Q}_j} [L_k(\phi, \omega)] \\ & \leq \limsup_{j \in \mathbb{N}} \mathbb{E}_{\omega \sim \mathbb{Q}_j} [\min\{L_k(\phi, \omega), M\}] + \limsup_{j \in \mathbb{N}} \mathbb{E}_{\omega \sim \mathbb{Q}_j} [(L_k(\phi, \omega) - M) \cdot \mathbb{1}_{\{L_k(\phi, \omega) > M\}}] \\ & \leq \mathbb{E}_{\omega \sim \mathbb{Q}} [\min\{L_k(\phi, \omega), M\}] + \limsup_{j \in \mathbb{N}} \mathbb{E}_{\omega \sim \mathbb{Q}_j} [(L_k(\phi, \omega) - M) \cdot \mathbb{1}_{\{L_k(\phi, \omega) > M\}}] \\ & \leq \mathbb{E}_{\omega \sim \mathbb{Q}} [\min\{L_k(\phi, \omega), M\}] + \sup_{\mathbb{Q} \in \mathbb{B}_{\bar{\rho}_k}(\mathbb{P}_k^\epsilon)} \mathbb{E}_{\omega \sim \mathbb{Q}} [L_k(\phi, \omega) \cdot \mathbb{1}_{\{L_k(\phi, \omega) > M\}}], \quad \forall k \in \mathbb{K}, \end{aligned} \quad (27)$$

where $M > 0$, the second inequality is due to Proposition 3.3 in [7] for the upper semicontinuous and bounded function $\min\{L_k(\phi, \omega), M\}$, and the third inequality follows because $\mathbb{Q}_j \in \mathbb{B}_{\bar{\rho}_k}(\mathbb{P}_k^\epsilon)$ for

any $j \in \mathbb{N}$ and $(L_k(\phi, \omega) - M) \cdot \mathbb{1}_{\{L_k(\phi, \omega) > M\}} \leq L_k(\phi, \omega) \cdot \mathbb{1}_{\{L_k(\phi, \omega) > M\}}$. Taking the limit as $M \rightarrow \infty$ on both sides of inequality (27), we obtain

$$\begin{aligned} & \limsup_{j \in \mathbb{N}} \mathbb{E}_{\omega \sim \mathbb{Q}_j} [L_k(\phi, \omega)] \\ & \leq \lim_{M \rightarrow \infty} \mathbb{E}_{\omega \sim \mathbb{Q}} [\min\{L_k(\phi, \omega), M\}] + \lim_{M \rightarrow \infty} \sup_{\mathbb{Q} \in \mathbb{B}_{\bar{\rho}_k}(\mathbb{P}_k^\epsilon)} \mathbb{E}_{\omega \sim \mathbb{Q}} [L_k(\phi, \omega) \cdot \mathbb{1}_{\{L_k(\phi, \omega) > M\}}] \\ & = \lim_{M \rightarrow \infty} \mathbb{E}_{\omega \sim \mathbb{Q}} [\min\{L_k(\phi, \omega), M\}] = \mathbb{E}_{\omega \sim \mathbb{Q}} [L_k(\phi, \omega)], \quad \forall k \in \mathbb{K}, \end{aligned}$$

where the first equality follows from Eq. (26) and the second equality follows from the monotone convergence theorem. This completes the proof. $\square$

B.2.3 Formal Proof of Theorem 1

Proof of Theorem 1. For Theorem 1(I), the problem (SDRHT) can be reformulated as

$$\inf_{\phi \in \mathcal{F}} \sup_{\mathbb{P}_k \in \mathbb{B}_{\bar{\rho}_k}(\widehat{\mathbb{P}}_k), \forall k \in \mathbb{K}} \sum_{k \in \mathbb{K}} \mathbb{E}_{\omega \sim \mathbb{P}_k} [L_k(\phi, \omega)].$$

Note that $\mathbb{B}_{\bar{\rho}_k}(\mathbb{P}_k^\epsilon)$ is weakly compact for any $\bar{\rho}_k \geq 0$ thanks to Proposition 1 and Lemma 1, which applies because of Assumption 1. In addition, $\mathbb{E}_{\omega \sim \mathbb{Q}} [L_k(\phi, \omega)]$ is non-negative and weakly upper semicontinuous in $\mathbb{Q} \in \mathbb{B}_{\bar{\rho}_k}(\mathbb{P}_k^\epsilon)$ for all $k \in \mathbb{K}$ and $\phi \in \mathcal{F}$, due to Definition 1 and Lemma 2. As $\Phi(\phi; \mathbb{P}_0, \mathbb{P}_1)$ is convex in ϕ and concave in distributions, Theorem 1(II) in [46], an extension of the lopsided minimax theorem [7, Theorem 5.15], ensures that

$$\text{Optimal Value of (SDRHT)} = \max_{\mathbb{P}_0 \in \mathbb{B}_{\bar{\rho}_0}(\widehat{\mathbb{P}}_0), \mathbb{P}_1 \in \mathbb{B}_{\bar{\rho}_1}(\widehat{\mathbb{P}}_1)} \inf_{\phi \in \mathcal{F}} \Phi(\phi; \mathbb{P}_0, \mathbb{P}_1).$$

For Theorem 1(II), the problem (soft-SDRHT) can be reformulated as

$$\inf_{\phi \in \mathcal{F}} \sup_{\mathbb{P}_0, \mathbb{P}_1 \in \mathcal{P}(\Omega)} \sum_{k \in \mathbb{K}} \mathbb{E}_{\omega \sim \mathbb{P}_k} [L_k(\phi, \omega)] - \lambda_k \mathcal{S}_\epsilon(\widehat{\mathbb{P}}_k, \mathbb{P}_k),$$

where $\lambda_k > 0$ for any $k \in \mathbb{K}$. Define

$$H(\phi, \mathbb{P}_0, \mathbb{P}_1) := \sum_{k \in \mathbb{K}} \mathbb{E}_{\omega \sim \mathbb{P}_k} [L_k(\phi, \omega)] - \lambda_k \mathcal{S}_\epsilon(\widehat{\mathbb{P}}_k, \mathbb{P}_k).$$

We next verify the level compactness of $H(\phi, \mathbb{P}_0, \mathbb{P}_1)$ required for the lopsided minimax theorem [7, Theorem 5.15]. By Eqs. (21), (23), and (24), for any fixed ϕ , we have

$$H(\phi, \mathbb{P}_0, \mathbb{P}_1) \leq \sum_{k \in \mathbb{K}} \left[C_k + \left(\frac{1}{\beta_k} - \lambda_k \epsilon \right) \cdot \inf_{\gamma \in \Gamma(\widehat{\mathbb{P}}_k, \mathbb{P}_k)} \mathbb{E}_{\mathbf{x} \sim \widehat{\mathbb{P}}_k} [\mathcal{D}_{\text{KL}}(\gamma_{\mathbf{x}} \| \mathcal{K}_{\mathbf{x}, \epsilon})] \right], \quad (28)$$

where $\beta_k > 0$ is finite, $C_k := \frac{1}{\beta_k} \log \mathbb{E}_{\omega \sim \mathbb{P}_k^\epsilon} [e^{\beta_k \cdot L_k(\phi, \omega)}] + \lambda_k \epsilon \cdot \mathbb{E}_{\mathbf{x} \sim \widehat{\mathbb{P}}_k} [\log \int_{\mathbb{R}^d} e^{-c(\mathbf{x}, \mathbf{z})/\epsilon} d\mathbf{z}]$ for any $k \in \mathbb{K}$ is finite by Assumptions 1(II) and 2(I), and is independent of $\mathbb{P}_0$ and $\mathbb{P}_1$.

By choosing $\beta_k \in (\frac{1}{\lambda_k \epsilon}, +\infty)$ for each $k \in \mathbb{K}$, when $\inf_{\gamma \in \Gamma(\widehat{\mathbb{P}}_k, \mathbb{P}_k)} \mathbb{E}_{\mathbf{x} \sim \widehat{\mathbb{P}}_k} [\mathcal{D}_{\text{KL}}(\gamma_{\mathbf{x}} \| \mathcal{K}_{\mathbf{x}, \epsilon})] \rightarrow \infty$, we have $H(\phi, \mathbb{P}_0, \mathbb{P}_1) \rightarrow -\infty$. This implies that for any $\alpha \in \mathbb{R}$, there exist finite radii r_0 and r_1 such that the upper-level set of $H(\phi, \mathbb{P}_0, \mathbb{P}_1)$ satisfies

$$\{(\mathbb{P}_0, \mathbb{P}_1) : H(\phi, \mathbb{P}_0, \mathbb{P}_1) \geq \alpha\} \subseteq \mathbb{B}_{r_0}(\widehat{\mathbb{P}}_0) \times \mathbb{B}_{r_1}(\widehat{\mathbb{P}}_1).$$

By Lemma 1, the set $\mathbb{B}_{r_0}(\widehat{\mathbb{P}}_0) \times \mathbb{B}_{r_1}(\widehat{\mathbb{P}}_1)$ is weakly compact. For each $k \in \mathbb{K}$, the Sinkhorn discrepancy is convex and weakly lower semicontinuous (see the proof of Lemma 1), and $\mathbb{E}_{\omega \sim \mathbb{P}_k} [L_k(\phi, \omega)]$ is weakly upper semicontinuous in $\mathbb{P}_k$ by Lemma 2. Hence, $H(\phi, \mathbb{P}_0, \mathbb{P}_1)$ is weakly upper semicontinuous and its upper-level set $\{(\mathbb{P}_0, \mathbb{P}_1) : H(\phi, \mathbb{P}_0, \mathbb{P}_1) \geq \alpha\}$ is weakly closed. Therefore, for every $\alpha \in \mathbb{R}$ and every fixed $\phi \in \mathcal{F}$, the upper-level set $\{(\mathbb{P}_0, \mathbb{P}_1) : H(\phi, \mathbb{P}_0, \mathbb{P}_1) \geq \alpha\}$, a closed subset of a weakly compact set, is also weakly compact. In addition, $-H(\phi, \mathbb{P}_0, \mathbb{P}_1)$ is convex, weakly lower semicontinuous, and thus weakly closed in $\mathbb{P}_0$ and $\mathbb{P}_1$.

Finally, $H(\cdot, \mathbb{P}_0, \mathbb{P}_1)$ is convex in ϕ , and Assumption 2(II) ensures that

$$\sup_{\mathbb{P}_0, \mathbb{P}_1} \inf_{\phi \in \mathcal{F}} H(\phi, \mathbb{P}_0, \mathbb{P}_1) > -\infty.$$

Therefore, as all conditions of the lopsided minimax theorem are satisfied, the strong duality of the problem (soft-SDRHT) holds, i.e.,

$$\text{Optimal Value of (soft-SDRHT)} = \sup_{\mathbb{P}_0, \mathbb{P}_1 \in \mathcal{P}(\Omega)} \inf_{\phi \in \mathcal{F}} \Phi(\phi; \mathbb{P}_0, \mathbb{P}_1) - \sum_{k \in \mathbb{K}} \lambda_k \mathcal{S}_\epsilon(\widehat{\mathbb{P}}_k, \mathbb{P}_k). \quad (29)$$

Further, since $\inf_{\phi \in \mathcal{F}} H(\phi, \mathbb{P}_0, \mathbb{P}_1)$ is weakly upper semicontinuous in $\mathbb{P}_0$ and $\mathbb{P}_1$, has a finite upper bound by Eq. (28) under proper β , and has weakly compact upper-level sets, the supremum in Eq. (29) is attained by Weierstrass' theorem, i.e.,

$$\text{Optimal Value of (soft-SDRHT)} = \max_{\mathbb{P}_0, \mathbb{P}_1 \in \mathcal{P}(\Omega)} \inf_{\phi \in \mathcal{F}} \Phi(\phi; \mathbb{P}_0, \mathbb{P}_1) - \sum_{k \in \mathbb{K}} \lambda_k \mathcal{S}_\epsilon(\widehat{\mathbb{P}}_k, \mathbb{P}_k).$$

This completes the proof. $\square$

B.3 Proof of Theorem 3

Proof of Theorem 3. As a preparation, when $\mathbb{D}$ has a finite diameter D , for any $v \in (0, 1)$, there exists a v -net $\mathbb{D}_v \subset \mathbb{D}$ containing J points such that $\forall \mathbf{x} \in \mathbb{D}, \exists \mathbf{x}_j \in \mathbb{D}_v, \|\mathbf{x} - \mathbf{x}_j\| \leq v$. According to [48], the size of the v -net is bounded by $J \leq \mathcal{O}((D/v)^d)$.

We prove Theorem 3(I) in two steps. According to Lemma 1 in [40], all continuous Lipschitz convex functions over convex compact sets can be approximated using the maximum of affine functions. Thus, in Step 1, we control the discretization error by adjusting the size J of the v -net. In Step 2, we prove that a HyCNN with maxout activation functions can exactly represent a maximum of affine functions and determine the quantitative relationship between J and the network hyperparameters (L, h) .

Step 1 of Theorem 3(I): For any L_f -Lipschitz function f , we approximate it by the maximum of supporting hyperplanes at points in the v -net. The supporting hyperplane at $(\mathbf{x}_j, f(\mathbf{x}_j))$ is defined as

$$s_j(\mathbf{x}) := \mathbf{g}_j^\top (\mathbf{x} - \mathbf{x}_j) + f(\mathbf{x}_j).$$

where $\mathbf{g}_j \in \partial f(\mathbf{x}_j)$. For any $\mathbf{x} \in \mathbb{D}$, let $\mathbf{x}_{j_0} \in \mathbb{D}_v$ be the nearest point to $\mathbf{x}$ in the v -net. The pointwise error satisfies

$$\begin{aligned} \sup_{\mathbf{x} \in \mathbb{D}} \left| f(\mathbf{x}) - \max_{j \in [J]} s_j(\mathbf{x}) \right| &= \sup_{\mathbf{x} \in \mathbb{D}} \left(f(\mathbf{x}) - \max_{j \in [J]} s_j(\mathbf{x}) \right) \\ &\leq \sup_{\mathbf{x} \in \mathbb{D}} \left(f(\mathbf{x}) - s_{j_0}(\mathbf{x}) \right) \\ &= \sup_{\mathbf{x} \in \mathbb{D}} \left(f(\mathbf{x}) - f(\mathbf{x}_{j_0}) - \mathbf{g}_{j_0}^\top (\mathbf{x} - \mathbf{x}_{j_0}) \right) \\ &\leq \sup_{\mathbf{x} \in \mathbb{D}} \left(|f(\mathbf{x}) - f(\mathbf{x}_{j_0})| + \|\mathbf{g}_{j_0}\| \cdot \|\mathbf{x} - \mathbf{x}_{j_0}\| \right) \\ &\leq 2L_f \|\mathbf{x} - \mathbf{x}_{j_0}\| \leq 2L_f v. \end{aligned}$$

Here, the first equality holds since f is convex, implying $s_j(\mathbf{x}) \leq f(\mathbf{x})$ for all j . The first inequality follows because $\max_{j \in [J]} s_j(\mathbf{x}) \geq s_{j_0}(\mathbf{x})$. The second inequality uses the triangle inequality and Cauchy-Schwarz inequality. The third inequality applies the L_f -Lipschitz continuity of f , which bounds both $|f(\mathbf{x}) - f(\mathbf{x}_{j_0})| \leq L_f \|\mathbf{x} - \mathbf{x}_{j_0}\|$ and $\|\mathbf{g}_{j_0}\| \leq L_f$. The last inequality holds as $\|\mathbf{x} - \mathbf{x}_{j_0}\| \leq v$. To ensure $\sup_{\mathbf{x} \in \mathbb{D}} |f(\mathbf{x}) - \max_{j \in [J]} s_j(\mathbf{x})| \leq \varepsilon$ for any $\varepsilon > 0$, we have $v \leq \min\{\frac{\varepsilon}{2L_f}, 1\}$. Taking $v = \varepsilon/2L_f$. By $J \leq \mathcal{O}((\frac{D}{v})^d)$, the number of points in the v -net is given by $J \leq \mathcal{O}((\frac{DL_f}{\varepsilon})^d)$.

Step 2 of Theorem 3(I): Here, we show that the maximum of J affine functions can be represented by a HyCNN with L layers and width h . As an intuitive example, we first consider a maximum of two affine functions and reformulate it as a HyCNN with two layers $\mathbf{y}_1$ and $\mathbf{y}_2$ (see Figure 2) by introducing an extra affine term:

$$\max \left\{ \mathbf{a}_1^\top \mathbf{x} + b_1, \mathbf{a}_2^\top \mathbf{x} + b_2 \right\} = \underbrace{(\mathbf{a}_3^\top \mathbf{x} + b_3)}_{\mathbf{y}_2} + \overbrace{\max \left\{ (\mathbf{a}_1 - \mathbf{a}_3)^\top \mathbf{x} + (b_1 - b_3), (\mathbf{a}_2 - \mathbf{a}_3)^\top \mathbf{x} + (b_2 - b_3) \right\}}^{\mathbf{y}_1}$$

We extend this intuition to a maximum of J affine functions $\max\{\mathbf{a}_1^\top \mathbf{x} + b_1, \dots, \mathbf{a}_J^\top \mathbf{x} + b_J\}$. We consider a simple case in which the width of each hidden layer is $h = 1$. In the following, we first construct a HyCNN with maxout activation functions that use both positive and negative weights, representing a maximum over affine functions. Then, all weights can be restricted to be nonnegative by the method in the proof of Theorem 1 in [40] without changing the width and depth of the network. Let $s_j := \mathbf{a}_j^\top \mathbf{x} + b_j$ for any $j \in [2J]$, extra affine terms $r_j := \mathbf{c}_j^\top \mathbf{x} + d_j$ for any

$j \in [2, \dots, J]$, $s'_j = s_j - r_j$ for any $j \in [2, \dots, J]$, and $\widehat{\sigma} : \mathbb{R}^2 \rightarrow \mathbb{R}$ be a maxout activation function. Then

$$\begin{aligned}
& \max\{s_1, \dots, s_J\} \\
&= \max\{\max\{s_1, \dots, s_{J-1}\}, s_J\} \\
&= r_J + \widehat{\sigma}(\max\{s_1, \dots, s_{J-1}\} - r_J, s'_J) \\
&= r_J + \widehat{\sigma}(\max\{\max\{s_1, \dots, s_{J-2}\}, s_{J-1}\} - r_J, s'_J) \\
&= r_J + \widehat{\sigma}(r_{J-1} + \widehat{\sigma}(\max\{s_1, \dots, s_{J-2}\} - r_{J-1}, s'_{J-1}) - r_J, s'_J) \\
&= \dots \\
&= r_J + \widehat{\sigma}(r_{J-1} + \widehat{\sigma}(r_{J-2} + \widehat{\sigma}(\dots + r_3 + \widehat{\sigma}(r_2 + \widehat{\sigma}(s_1 - r_2, s'_2) - r_3, s'_3) - \dots) - r_{J-2}, s'_{J-1}) - r_J, s'_J).
\end{aligned} \tag{30}$$

The last equation describes a J -layer HyCNN with width $h = 1$, where the layers are

$$\begin{aligned}
y_1 &= \widehat{\sigma}(s_1 - r_2, s_2 - r_2), \\
y_2 &= \widehat{\sigma}(1 \cdot y_1 + r_2 - r_3, 0 \cdot y_1 + s_3 - r_3), \\
&\dots \\
y_j &= \widehat{\sigma}(1 \cdot y_{j-1} + r_j - r_{j+1}, 0 \cdot y_{j-1} + s_{j+1} - r_{j+1}), \\
&\dots \\
y_J &= r_J + y_{J-1}.
\end{aligned}$$

To ensure the coefficients of previous layers in this neural network are nonnegative, following [40], we expand the original input $\mathbf{x} \in \mathbb{R}^d$ in each layer to $\widehat{\mathbf{x}} \in \mathbb{R}^{2d}$ that includes $\mathbf{x}$ and $-\mathbf{x}$, i.e., $\widehat{\mathbf{x}} = [\mathbf{x}^\top, -\mathbf{x}^\top]^\top$. When all coefficients of $\mathbf{x}$ are nonnegative, the coefficients of $-\mathbf{x}$ are 0. When $h_j < 0$, we set the coefficient of $\widehat{\mathbf{x}}_{j+d}$ to $-h_j$ and that of $\mathbf{x}_j$ to 0. Therefore, without loss of generality, we can limit all weights between layers to be nonnegative, thereby ensuring that the neural network is input-convex.

The construction in Eq. (30) represents $\max_{j \in [J]} s_j(\mathbf{x})$ exactly by a HyCNN with width $h = 1$ and depth $L = \mathcal{O}(J)$. By Step 1, choosing $h = 1$ and $L = \mathcal{O}(J)$ gives

$$L \cdot h = \mathcal{O}\left(\left(\frac{DL_f}{\varepsilon}\right)^d\right),$$

which proves Theorem 3(I).

For Theorem 3(II), we first focus on the property of the smooth maxout activation function. Without loss of generality, suppose $a_1 \geq a_2$, we have $\bar{\sigma}(a_1, a_2) = a_1$ and

$$\bar{\sigma}_\tau(a_1, a_2) = \tau \log\left(e^{a_1/\tau} (1 + e^{(a_2 - a_1)/\tau})\right) = a_1 + \tau \log\left(1 + e^{(a_2 - a_1)/\tau}\right) \leq a_1 + \tau \log 2,$$

where $\tau > 0$ and the last inequality is due to $e^{(a_2 - a_1)/\tau} \in (0, 1]$. This is equivalent to

$$\bar{\sigma}(a_1, a_2) < \bar{\sigma}_\tau(a_1, a_2) \leq \bar{\sigma}(a_1, a_2) + \tau \log 2.$$

Consequently, for a HyCNN and its smooth version, we have

$$\sup_{\mathbf{x} \in \mathbb{D}} |\widehat{\text{HyCNN}}^\tau(\mathbf{x}) - \widehat{\text{HyCNN}}(\mathbf{x})| \leq C \cdot \tau \log 2$$

where $C > 0$ is a constant depending on the number of maxout activation functions and layers in the network. By Theorem 3(I), for any $\varepsilon > 0$, there exists a HyCNN such that

$$\sup_{x \in \mathbb{D}} |\widehat{\text{HyCNN}}(x) - f(x)| \leq \frac{\varepsilon}{2}.$$

By selecting a sufficiently small τ , e.g., $\tau \leq \frac{\varepsilon}{2C \log 2}$, we have

$$\sup_{x \in \mathbb{D}} |\widehat{\text{HyCNN}}^\tau(x) - f(x)| \leq \sup_{x \in \mathbb{D}} |\widehat{\text{HyCNN}}^\tau(x) - \widehat{\text{HyCNN}}(x)| + \sup_{x \in \mathbb{D}} |\widehat{\text{HyCNN}}(x) - f(x)| \leq \varepsilon.$$

Note that the smooth HyCNN and standard HyCNN share the same order of network architecture complexity. Hence, the hyperparameters of smooth HyCNN can also be chosen as $L \cdot h = \mathcal{O}((\frac{DL_f}{\varepsilon})^d)$ to guarantee the ε -approximation. This completes the proof of Theorem 3(II).

For Theorem 3(III), we select a sequence of smooth HyCNNs on a compact domain $[-n, n]^d$ with sufficiently small τ and each satisfying $L_n \cdot h_n = \mathcal{O}((\frac{DL_f}{\varepsilon})^d)$. By setting $\varepsilon_n = 1/n$, according to Theorem 3(II), we have $\sup_{x \in \mathbb{D}} |\widehat{\text{HyCNN}}_{\varepsilon_n}^\tau(x) - f(x)| \leq \frac{1}{n}$ where the right-hand-side term converges to 0 as $n \rightarrow \infty$. Hence, we can prove $\widehat{\text{HyCNN}}_n^\tau$ converges to f uniformly on $\mathbb{D}$. This completes the proof. $\square$

B.4 Proof of Theorem 4

Proof of Theorem 4. For Theorem 4(I), one can see Theorem 2 in [34] for a detailed proof.

For Theorem 4(II), let the function f be M -Lipschitz smooth, and then we have

$$f(y) - f(x) \leq \nabla f(x)^\top (y - x) + \frac{M}{2} \cdot \|y - x\|^2. \quad (31)$$

For any $\widehat{\text{HyCNN}}_n^\tau$ in the sequence $\{\widehat{\text{HyCNN}}_n^\tau\}_{n=1}^\infty$ with hyperparameters satisfying $L_n \cdot h_n = \mathcal{O}(\varepsilon_n^{-d})$ for any $\varepsilon_n > 0$, we have

$$\widehat{\text{HyCNN}}_n^\tau(y) - \widehat{\text{HyCNN}}_n^\tau(x) \geq \nabla \widehat{\text{HyCNN}}_n^\tau(x)^\top (y - x), \quad \forall x, y \in \mathbb{R}^d. \quad (32)$$

By Eqs. (31) and (32), we have

$$\left(\nabla \widehat{\text{HyCNN}}_n^\tau(x) - \nabla f(x) \right)^\top (y - x) \leq \left(\widehat{\text{HyCNN}}_n^\tau(y) - f(y) \right) - \left(\widehat{\text{HyCNN}}_n^\tau(x) - f(x) \right) + \frac{M}{2} \cdot \|y - x\|^2.$$

Let y be obtained by moving x along the direction $\nabla \widehat{\text{HyCNN}}_n^\tau(x) - \nabla f(x)$ with a proper step length η , i.e., $y - x = \eta \frac{\nabla \widehat{\text{HyCNN}}_n^\tau(x) - \nabla f(x)}{\|\nabla \widehat{\text{HyCNN}}_n^\tau(x) - \nabla f(x)\|}$ and $y \in \mathbb{D}$. Then, for the point-wise supremum value of $\nabla \widehat{\text{HyCNN}}_n^\tau(x) - \nabla f(x)$ on compact set $\mathbb{D} \subset \mathbb{R}^d$, we have

$$\begin{aligned} \sup_{x \in \mathbb{D}} \left\| \nabla \widehat{\text{HyCNN}}_n^\tau(x) - \nabla f(x) \right\| &\leq \frac{1}{\eta} \cdot \sup_{x \in \mathbb{D}} \left(\widehat{\text{HyCNN}}_n^\tau(y) - f(y) - \left(\widehat{\text{HyCNN}}_n^\tau(x) - f(x) \right) \right) + \frac{M}{2\eta} \cdot \|y - x\|^2 \\ &\leq \frac{2\varepsilon_n}{\eta} + \frac{M\eta}{2} \leq 2\sqrt{M\varepsilon_n} = \mathcal{O}(\sqrt{\varepsilon_n}), \end{aligned}$$

where the second inequality is due to Theorem 3(II) and $\|y - x\|^2 = \eta^2$, the last inequality is due to the basic inequality, and the last equality is due to $L_n \cdot h_n = \mathcal{O}(\varepsilon_n^{-d})$ for any $\varepsilon_n > 0$. Hence, the gradient sequence $\nabla \widehat{\text{HyCNN}}_n$ converges to ∇f uniformly at a rate of $\mathcal{O}(\sqrt{\varepsilon_n})$ when function f is Lipschitz smooth. This completes the proof. $\square$

B.5 Proof of Theorem 5

Proof of Theorem 5. Since μ is absolutely continuous and both μ and ν have finite second moments, Brenier's theorem yields a finite convex function $f : \mathbb{R}^d \rightarrow \mathbb{R}$ such that

$$(\nabla f)_\# \mu = \nu.$$

Since every finite convex function on $\mathbb{R}^d$ is locally Lipschitz, for every $n \in \mathbb{N}$, f is Lipschitz on a neighborhood of the compact set

$$\mathbb{D}_n = [-n, n]^d.$$

By Theorem 3 (II), there exists a smooth HyCNN $\widehat{\text{HyCNN}}_n^{\tau_n}$ satisfying

$$\sup_{x \in \mathbb{D}_n} \left| \widehat{\text{HyCNN}}_n^{\tau_n}(x) - f(x) \right| \leq \frac{1}{n}.$$

For any fixed compact set $\mathbb{K} \subset \mathbb{R}^d$, there exists N such that $\mathbb{K} \subseteq \mathbb{D}_n$ for every $n \geq N$. Hence,

$$\sup_{x \in \mathbb{K}} \left| \widehat{\text{HyCNN}}_n^{\tau_n}(x) - f(x) \right| \leq \frac{1}{n} \rightarrow 0.$$

Thus, $\widehat{\text{HyCNN}}_n^{\tau_n} \rightarrow f$ locally uniformly, and in particular pointwise on $\mathbb{R}^d$.

By Theorem 4 (I),

$$\nabla \widehat{\text{HyCNN}}_n^{\tau_n}(x) \rightarrow \nabla f(x)$$

at almost every $x \in \mathbb{R}^d$. Since μ is absolutely continuous with respect to the Lebesgue measure, this implies the weak convergence of the pushforward measure holds μ -almost everywhere, i.e., for $X \sim \mu$,

$$\nabla \widehat{\text{HyCNN}}_n^{\tau_n}(X) \rightarrow \nabla f(X) \quad \text{almost surely.}$$

Consequently,

$$(\nabla \widehat{\text{HyCNN}}_n^{\tau_n})_\# \mu \Rightarrow (\nabla f)_\# \mu = \nu,$$

where $\Rightarrow$ denotes weak convergence of probability measures on $\mathbb{R}^d$. This completes the proof. $\square$

C Technical Note for Training and Testing HyCNN

In this section, we provide details to compute the gradients of the problem (7) in Algorithm 1.

C.1 Gradient of the Risk Function Estimation

In this subsection, we take the $k = 0$ term in Eq. (6) as an example. For brevity, we omit the superscript (i) . Denote the parameters of the neural network HyCNN_k as θ_k for each $k \in \mathbb{K}$. Since our approximated convex function $\widetilde{\varphi}_k$ is strongly convex, we have $\det \nabla^2 \widetilde{\varphi}_k(z) > 0$ for any k and z .

(I) **Gradient with respect to θ_0 .**

For brevity, let $h(\theta_0, \theta_1) := \exp\left[\frac{1}{2}(\log p_1(\omega_0) - \log p_0(\omega_0))\right]$, which is the same as $h_0(\omega_0, \theta_0, \theta_1)$ in Algorithm 3. Then, we have

$$\frac{\partial h(\theta_0, \theta_1)}{\partial \theta_0} = \frac{h(\theta_0, \theta_1)}{2} \cdot \left[\frac{\partial \log p_1(\omega_0)}{\partial \theta_0} - \frac{\partial \log p_0(\omega_0)}{\partial \theta_0} \right]. \quad (33)$$

For the second item on the right-hand side of Eq. (33), based on Eq. (8), we have

$$\frac{\partial \log p_0(\omega_0)}{\partial \theta_0} = \frac{\partial}{\partial \theta_0} \log p_0^\epsilon(z_0) - \frac{\partial}{\partial \theta_0} \log \det \nabla^2 \tilde{\varphi}_0(z_0) = -\frac{\partial}{\partial \theta_0} \log \det \nabla^2 \tilde{\varphi}_0(z_0), \quad (34)$$

where $z_0 = \nabla \varphi_0^{-1}(\omega_0)$, the first equality is due to the strong convexity of $\tilde{\varphi}_0$, and the second equality is due to $\frac{\partial}{\partial \theta_0} \log p_0^\epsilon(z_0) = 0$. The log-determinant-Hessian in Eq. (34) can be efficiently computed by Algorithm 2.

For the first item on the right-hand side of Eq. (33), based on the chain rule, we have

$$\frac{\partial \log p_1(\omega_0)}{\partial \theta_0} = \left(\frac{\partial \omega_0}{\partial \theta_0} \right)^\top \left(\frac{\partial \tilde{z}_0}{\partial \omega_0} \right)^\top \frac{\partial \log p_1(\omega_0)}{\partial \tilde{z}_0}. \quad (35)$$

Let vector

$$\mathbf{g} := \frac{\partial \log p_1(\omega_0)}{\partial \tilde{z}_0} = \frac{\partial \left[\log p_1^\epsilon(\tilde{z}_0) - \log \det \nabla^2 \tilde{\varphi}_1(\tilde{z}_0) \right]}{\partial \tilde{z}_0},$$

and this derivative can be computed by automatic differentiation. Let $\tilde{z}_0 = \nabla \tilde{\varphi}_1^{-1}(\omega_0)$. As $\frac{\partial \tilde{z}_0}{\partial \omega_0} = [\nabla^2 \tilde{\varphi}_1(\tilde{z}_0)]^{-1}$, Eq.(35) can be reformulated as

$$\frac{\partial \log p_1(\omega_0)}{\partial \theta_0} = \left(\frac{\partial \omega_0}{\partial \theta_0} \right)^\top [\nabla^2 \tilde{\varphi}_1(\tilde{z}_0)]^{-1} \mathbf{g}, \quad (36)$$

where $[\nabla^2 \tilde{\varphi}_1(\tilde{z}_0)]^{-1} \mathbf{g}$ can be efficiently solved as the minimizer of a quadratic optimization problem similar to (15) without explicitly computing the inverse of the Hessian matrix. Then, $\left(\frac{\partial \omega_0}{\partial \theta_0} \right)^\top [\nabla^2 \tilde{\varphi}_1(\tilde{z}_0)]^{-1} \mathbf{g}$ can be computed by vector-Jacobian product in an automatic differentiation framework. In summary, the gradient of $h(\theta_0, \theta_1)$ with respect to θ_0 is given by

$$\frac{\partial h(\theta_0, \theta_1)}{\partial \theta_0} = \frac{h(\theta_0, \theta_1)}{2} \cdot \left[\left(\frac{\partial \omega_0}{\partial \theta_0} \right)^\top [\nabla^2 \tilde{\varphi}_1(\tilde{z}_0)]^{-1} \mathbf{g} + \frac{\partial}{\partial \theta_0} \log \det \nabla^2 \tilde{\varphi}_0(z_0) \right].$$

(II) **Gradient with respect to θ_1 .**

Similar to Eq. (33), we have

$$\frac{\partial h(\theta_0, \theta_1)}{\partial \theta_1} = \frac{h(\theta_0, \theta_1)}{2} \cdot \left[\frac{\partial \log p_1(\omega_0)}{\partial \theta_1} - \frac{\partial \log p_0(\omega_0)}{\partial \theta_1} \right] = \frac{h(\theta_0, \theta_1)}{2} \cdot \frac{\partial \log p_1(\omega_0)}{\partial \theta_1}. \quad (37)$$

Since $\tilde{z}_0 = \nabla \tilde{\varphi}_1^{-1}(\omega_0)$, we have

$$\begin{aligned} \frac{\partial \log p_1(\omega_0)}{\partial \theta_1} &= \frac{\partial}{\partial \theta_1} \left[\log p_1^\epsilon(\tilde{z}_0) - \log \det \nabla^2 \tilde{\varphi}_1(\tilde{z}_0) \right] \\ &= \left(\frac{\partial \tilde{z}_0}{\partial \theta_1} \right)^\top \frac{\partial \log p_1^\epsilon(\tilde{z}_0) - \log \det \nabla^2 \tilde{\varphi}_1(\tilde{z}_0)}{\partial \tilde{z}_0} - \frac{\partial}{\partial \theta_1} \log \det \nabla^2 \tilde{\varphi}_1(z) \Big|_{z=\tilde{z}_0}, \end{aligned} \quad (38)$$

where the second equality is based on the total derivative. To evaluate the Jacobian $\frac{\partial \bar{z}_0}{\partial \theta_1}$ in the first term, we take the total derivative with respect to θ_1 on both sides of $\nabla \tilde{\varphi}_1(\bar{z}_0) = \omega_0$:

$$\nabla^2 \tilde{\varphi}_1(\bar{z}_0) \cdot \frac{\partial \bar{z}_0}{\partial \theta_1} + \frac{\partial \nabla \tilde{\varphi}_1(z)}{\partial \theta_1} \Big|_{z=\bar{z}_0} = \mathbf{0},$$

and it follows that

$$\left(\frac{\partial \bar{z}_0}{\partial \theta_1}\right)^\top = -\left[\nabla^2 \tilde{\varphi}_1(\bar{z}_0)\right]^{-1} \frac{\partial \nabla \tilde{\varphi}_1(\bar{z}_0)}{\partial \theta_1}.$$

Therefore, Eq. (38) can be reformulated as

$$\frac{\partial \log p_1(\omega_0)}{\partial \theta_1} = -\left(\frac{\partial \nabla \tilde{\varphi}_1(\bar{z}_0)}{\partial \theta_1}\right)^\top \left[\nabla^2 \tilde{\varphi}_1(\bar{z}_0)\right]^{-1} \mathbf{g} - \frac{\partial}{\partial \theta_1} \log \det \nabla^2 \tilde{\varphi}_1(\bar{z}_0),$$

where the first term can be calculated in the same way as Eq. (36) and the second term can be computed by Algorithm 2 as $\bar{z}_0$ is a constant here, thanks to the total derivative. In summary, the gradient of $h(\theta_0, \theta_1)$ with respect to θ_1 is given by

$$\frac{\partial h(\theta_0, \theta_1)}{\partial \theta_1} = -\frac{h(\theta_0, \theta_1)}{2} \cdot \left[\left(\frac{\partial \nabla \tilde{\varphi}_1(\bar{z}_0)}{\partial \theta_1}\right)^\top \left[\nabla^2 \tilde{\varphi}_1(\bar{z}_0)\right]^{-1} \mathbf{g} + \frac{\partial}{\partial \theta_1} \log \det \nabla^2 \tilde{\varphi}_1(\bar{z}_0)\right].$$

C.2 Gradient of the Sinkhorn Discrepancy Estimation

When the potential f is optimally solved, the total derivative of our Sinkhorn discrepancy estimator (18) with respect to the HyCNN parameters is given by

$$\frac{\partial \tilde{\mathcal{S}}_{\epsilon, (n_k, N)}(\hat{\mathbb{P}}_k, \mathbb{P}_k)}{\partial \theta_k} = \frac{1}{N} \sum_{i=1}^N \frac{\partial g_{\epsilon, k}^c(\omega_k^{(i)})}{\partial \theta_k} \Big|_{f=f_{\epsilon, k, (n_k, N)}} = \frac{1}{N} \sum_{i=1}^N \left(\frac{\partial \omega_k^{(i)}}{\partial \theta_k}\right)^\top \cdot \frac{\partial g_{\epsilon, k}^c(\omega_k^{(i)})}{\partial \omega_k^{(i)}} \Big|_{f=f_{\epsilon, k, (n_k, N)}}. \quad (39)$$

where the first equality is due to the envelope theorem, $\frac{\partial \omega_k^{(i)}}{\partial \theta_k}$ is the Jacobian of the generated sample with respect to the HyCNN parameters, and the gradient of $g_{\epsilon, k}^c$ with respect to $\omega_k^{(i)}$ can be explicitly computed based on Eq. (17).

C.3 Testing

In the testing phase, for a testing sample $\omega \in \Omega$, we first trace its position in distributions $\mathbb{P}_0^\epsilon$ and $\mathbb{P}_1^\epsilon$ by $\bar{z}_0 = \nabla \tilde{\varphi}_0^{-1}(\omega)$ and $\bar{z}_1 = \nabla \tilde{\varphi}_1^{-1}(\omega)$, respectively. Then, we evaluate the log-density of least-favorable distributions $\mathbb{P}_0$ and $\mathbb{P}_1$ at ω by

$$\log p_k(\omega) = \log p_k^\epsilon(\bar{z}_k) - \log \det \nabla^2 \tilde{\varphi}_k(\bar{z}_k), \quad \forall k \in \mathbb{K}.$$

Finally, we make decisions by the selected optimal detector $\phi_{\text{opt}} = \frac{1}{2}(\log p_0(\omega) - \log p_1(\omega))$.

D Wasserstein Distributionally Robust Hypothesis Testing Model

In this section, we provide the model of the WDRO method that is used as a baseline in our experimental part. The WDRO method is the soft-constraint version of the Wasserstein distributionally robust hypothesis testing model proposed by [17] (assuming $n_0 = n_1 = n$):

$$\begin{aligned}
& \max_{\mathbf{p}_0, \mathbf{p}_1 \in \mathbb{R}_+^n, \gamma_0, \gamma_1 \in \mathbb{R}_+^{n \times n}} \sum_{i \in [n]} \min\{\mathbf{p}_0^i, \mathbf{p}_1^i\} - \sum_{k \in \mathbb{K}} \lambda_k \cdot \sum_{i \in [n]} \sum_{j \in [n]} \gamma_{k,i,j} \cdot c(\mathbf{x}_k^{(i)}, \mathbf{x}_k^{(j)}) \\
& \text{s.t.} \quad \sum_{j \in [n]} \gamma_{k,i,j} = \widehat{\mathbb{P}}_k^i \quad \forall i \in [n], k \in \mathbb{K}, \\
& \quad \sum_{i \in [n]} \gamma_{k,i,j} = \mathbf{p}_k^j \quad \forall j \in [n], k \in \mathbb{K},
\end{aligned} \tag{40}$$

where $\lambda_k \geq 0$ is a penalty parameter for each $k \in \mathbb{K}$, $\gamma_{k,i,j}$ represents the ij -th entry of γ_k , and the i -th entry of $\mathbf{p}_k$ (respectively, $\widehat{\mathbb{P}}_k$) is specified by $\mathbf{p}_k^i$ (respectively, $\widehat{\mathbb{P}}_k^i$). This model assumes that the samples in the testing set have the same support as the empirical distribution. To extend the resulting least-favorable distribution to the entire space, a kernel smoothing method is introduced. Let $\mathbf{p}_k^*$ be an optimal solution of (40). This method convolves the discrete distribution with a kernel function $G_h : \mathbb{R}^d \rightarrow \mathbb{R}$ parameterized by a bandwidth $h > 0$:

$$\mathbf{p}_k^h(\boldsymbol{\omega}) = \sum_{i=1}^n \mathbf{p}_k^*(\mathbf{x}_k^{(i)}) \cdot G_h(\boldsymbol{\omega} - \mathbf{x}_k^{(i)}), \quad \forall \boldsymbol{\omega} \in \Omega, k \in \mathbb{K},$$

where we choose the Gaussian kernel function $G_h(\mathbf{x}) = \exp(-\|\mathbf{x}\|^2/2h^2)$. Note that the number of decision variables of this linear program is $\mathcal{O}(n_0 \cdot n_1)$. Thus, this finite-dimensional convex program is not sensitive to dimension d , but is sensitive to the number of training data points of the empirical distribution.