

CyclOT: Learning Quadratic Optimal Transport Maps via Synchronized Forward–Backward Interpolants

Shizhou Xu ^{*} ¹, Jiachen Liu ², Shih-Hsin Wang [†] ³, Stefan Broecker ⁴, Yuhao Huang ⁵,
Bao Wang ⁵, and Thomas Strohmer ²

¹SLAC National Accelerator Laboratory, Stanford University, Menlo Park, CA 94025, USA

²Department of Mathematics, University of California, Davis, Davis, CA 95616, USA

³Department of Computer Science and Information Engineering, National Taiwan University, Taipei 10617, Taiwan

⁴Department of Computer Science, University of California, Davis, Davis, CA 95616, USA

⁵Department of Mathematics and Scientific Computing and Imaging Institute, University of Utah, Salt Lake City, UT 84112, USA

Abstract

We study the recovery of forward and reverse quadratic optimal-transport maps from unpaired samples in high dimensions. We introduce a bidirectional neural framework in which the learned maps induce forward and backward displacement interpolants, while the training objective combines bidirectional quadratic action, discriminator-restricted Jensen-Shannon endpoint objectives, and two-sided cycle consistency. The construction requires neither precomputed sample pairings nor an explicit convex-potential parameterization. For absolutely continuous probability measures supported on a compact convex set, and under the stated generator-approximation, discriminator-richness, and minimizer-attainment conditions, we prove a population recovery theorem: for every prescribed accuracy, the sum of the corresponding L^2 errors between any global minimizer and the forward and reverse quadratic Brenier maps is below that accuracy, provided the discriminator level is sufficiently large and the annealing action weight becomes sufficiently small. Moreover, the cycle loss is bounded above by $\lambda W_2^2(\mu_0, \mu_1)$. Complementary results quantify approximate invertibility and show that exact endpoint Jensen-Shannon divergence and cycle consistency control missing target mass and many-to-one map collapse, respectively. Experiments on Swiss roll, MNIST, CelebA, single-cell perturbation data, and chest X-ray images evaluate endpoint fidelity, transport cost, inverse consistency, and the geometry of the induced interpolations.

Keywords: optimal transport, Brenier maps, displacement interpolation, transport annealing, cycle consistency.

^{*}Corresponding author: shzxu@stanford.edu. This work was performed while the author was affiliated with the Department of Mathematics, University of California, Davis.

[†]This work was performed while the author was affiliated with the Department of Mathematics, University of Utah.

Contents

1	Introduction	4
1.1	Our contributions	5
1.2	Related work	6
1.3	Organization	7
2	Problem Setup & Notation	7
2.1	Quadratic optimal transport	7
2.2	Displacement interpolation	8
2.3	Bidirectional learning problem	9
3	Proposed Framework	9
3.1	Forward-backward map framework	10
3.2	Neural population objective	10
3.3	Practical optimization algorithm	11
4	Global Recovery of the Quadratic Brenier Maps	13
4.1	Neural approximation ingredients	13
4.1.1	Generator approximation & range preservation	13
4.1.2	Joint map & cycle approximation	14
4.1.3	Discriminator uniform approximation	15
4.1.4	Fixed-level discriminator compactness & continuity	16
4.2	Recovery of the forward-backward Brenier pair	19
4.3	Discriminator refinement & action-weight annealing	24
5	Structural Roles of Endpoint Matching, Cycle Consistency, & Action	24
5.1	Feasible-class reduction via cycle consistency	26
5.2	Endpoint mode coverage and suppression of map-level collapse	27
5.3	Local interaction among action, neural JSD, & cycle consistency	29
6	Experiments	31
6.1	Experimental protocol & interpretation of metrics	31
6.2	Swiss roll: geometric fidelity in low dimensions	33
6.3	MNIST: transport from digits 0–4 to digits 5–9	34
6.4	CelebA: bidirectional transport on natural images	36
6.5	Single-cell data: transport on scientific manifolds	37
6.6	Chest X-ray translation: transport on medical images	38
6.7	Takeaway	40
A	Background & Preliminaries	46
A.1	Flow matching & rectified flow	46
A.2	Independence constraints & Wasserstein barycenters.	46
B	Proofs	47
B.1	Proof of Lemma 4.1	47
B.2	Proof of Appendix 5.1	48

C	Evaluation Metrics & Experimental Details	49
C.1	Evaluation metrics	49
C.1.1	Transport cost evaluation	49
C.1.2	Marginal matching quality	49
C.1.3	Efficiency	51
D	Additional Experimental Results	52
D.1	MNIST	52
D.2	CelebA	53
D.3	Pneumonia	55
D.4	Ablation on cycle-consistency loss	55

1 Introduction

Learning meaningful transformations between unpaired high-dimensional distributions is a fundamental challenge in modern machine learning. Such problems arise in image-to-image translation [1], domain adaptation [2], single-cell perturbation modeling [3], and learning straight flow paths for flow-based generative models [4], where ground-truth correspondences are inherently unobserved and matching endpoint marginals alone is insufficient. The goal is instead to recover a structurally meaningful map that reflects the underlying geometry or dynamics. Optimal transport (OT) provides a principled framework for this task by selecting, among all maps matching the source and target distributions, one that minimizes the expected transport cost [5, 6, 7]: under quadratic cost, the Brenier map is the gradient of a convex potential [8], and the induced McCann displacement interpolation defines a Wasserstein geodesic between the two distributions [9]. Equivalently, the Benamou–Brenier dynamic formulation characterizes this path as the minimum-kinetic-energy solution of a continuity equation [10].

Despite this appealing characterization, learning OT maps remains difficult in high dimensions. Classical computational OT methods, including entropic Sinkhorn iterations [11, 7], are primarily discrete, can be expensive at large sample sizes, and do not directly provide parametric maps for out-of-sample generalization. Recent flow-matching methods improve scalability by using prescribed couplings between source and target samples [4]. In particular, minibatch OT couplings are used to straighten flow trajectories and improve few-step generation [12, 13]. However, these couplings are batch-dependent approximations of the population OT plan and may introduce both computational overhead and statistical bias [14].

Neural OT methods instead seek to learn transport maps or potentials directly from samples. Dual and minimax formulations, including Wasserstein adversarial objectives and neural OT methods [15, 16], are principled but can be sensitive to optimization and provide limited direct supervision of the intermediate transport path. Brenier-inspired approaches [17, 18, 19] represent the quadratic OT map as the gradient of a convex potential, often parameterized by input convex neural networks (ICNNs) [8, 20]. While this preserves the convex-potential structure required by Brenier’s theorem, ICNN-based parameterizations can be restrictive and difficult to optimize compared with standard neural architectures. Optimal flow matching [21] is closely related in its goal of learning straight OT trajectories, but still belongs to the broader class of methods whose practical performance depends on surrogate coupling or potential-learning mechanisms. As a result, existing methods often trade scalability, architectural flexibility, optimization stability, and fidelity for the population OT solution.

Very recent work further highlights the importance of coupling design in flow-based generative modeling. Semi-discrete OT coupling methods replace minibatch OT with dataset-level semi-discrete formulations to improve scalability and trajectory straightness [22, 23], while hierarchical or dependent-coupling approaches modify the coupling structure to improve few-step or multistage generation [24, 25]. These methods reinforce the importance of coupling structure, but they still construct couplings or interpolants through external pairing mechanisms. In contrast, our method, CyclOT, couples the induced intermediate interpolants through jointly learned forward and backward maps, with cycle consistency promoting approximate invertibility and coherence between the two transport directions, without requiring precomputed pairings or minibatch/semi-discrete OT subproblems during training.

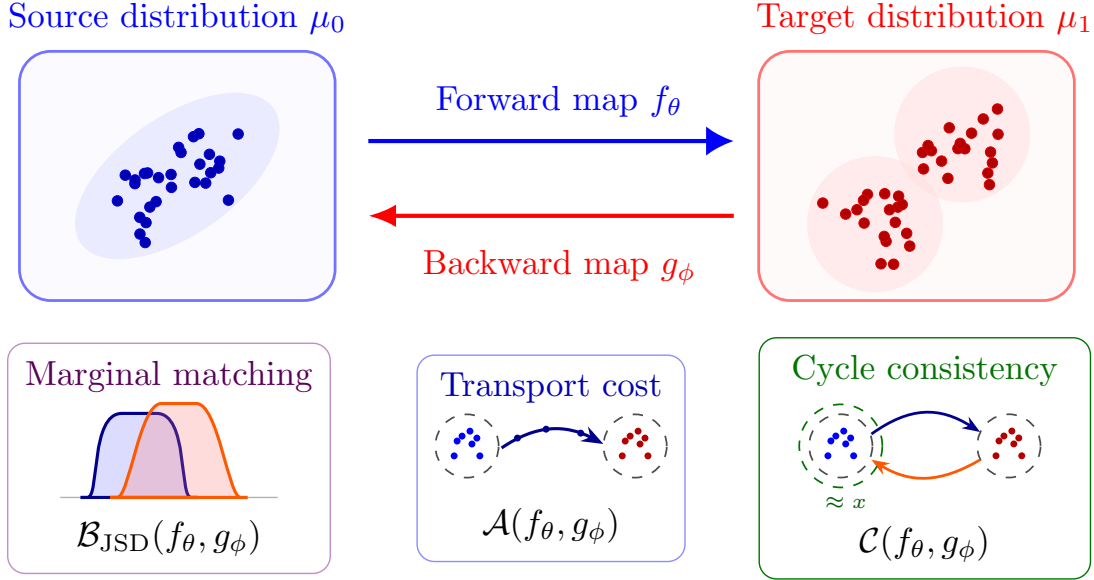


Figure 1: Overview of our proposed **target-free OT learning via synchronized forward-backward interpolants**. The forward map f_θ transports μ_0 from left to right, while the backward map g_ϕ transports μ_1 from right to left; Both f_θ and g_ϕ are neural network-approximated maps. The forward and backward induced marginals are matched at the two endpoints, providing marginal supervision for the transport maps.

1.1 Our contributions

We propose a new framework, **CyclOT**, for learning the quadratic OT map directly from unpaired samples, without iterative Reflow-style retraining [4], minibatch or semi-discrete OT coupling solvers [12, 13, 22, 23], or restrictive convex-potential parameterizations [17, 18, 19, 20]. Our method learns two maps simultaneously: a forward transport map from source to target and a backward map from target to source. These maps induce forward and backward interpolants evolving from opposite endpoints, while cycle consistency couples the two transport directions and promotes coherent intermediate trajectories. The resulting objective can be approximated by two neural networks and optimized using stochastic optimization. At the theoretical level, our objective is motivated by an L^2 projection view of Wasserstein barycenters under independence constraints [26]. This perspective suggests that enforcing independence-like intermediate agreement between map-induced interpolants can identify the quadratic OT structure without explicitly solving for a coupling.

Our main contributions are:

1. **A target-free learning framework for OT maps.** We introduce synchronized forward-backward interpolants that impose structural consistency along the transport path without requiring paired samples, precomputed couplings, or externally specified trajectories. See Section 3 for details.
2. **Scalable training with standard neural architectures.** Our method, CyclOT, avoids iterative Reflow-style retraining [4], minibatch OT approximations [12, 13], and ICNNs [20], enabling efficient optimization with flexible expressive architectures such as ResNets or Transformers. See Section 3.2 for details.

3. **Provable recovery of the quadratic OT map.** Under suitable assumptions, we show that minimizing our objective recovers the population quadratic OT map and its inverse. See Section 4.2 for details.
4. **Empirical validation across generative and scientific benchmarks.** Across synthetic, image, and scientific datasets, our approach yields stable training, accurate one-step generation, and meaningful interpolations competitive with representative OT-based and flow-based generative methods. See Section 6 for details.

1.2 Related work

Neural Optimal Transport and Direct Map Learning. Classical OT provides a rigorous variational formulation for matching probability measures [5, 6, 7], and the quadratic-cost case admits the Brenier-map structure [8]. However, computing OT maps from samples in high dimension remains challenging. Entropic solvers such as Sinkhorn iterations [11] are effective for discrete OT but do not directly yield parametric out-of-sample maps. Neural OT methods address this by parameterizing transport maps, costs, or Kantorovich potentials with neural networks [16]. A prominent Brenier-inspired line represents maps as gradients of convex potentials, often implemented using ICNNs [17, 18, 19, 20]. While theoretically appealing, such convex-potential parameterizations can restrict architectural flexibility and complicate optimization. CyclOT instead learns forward and backward maps using standard neural network architectures and enforces OT structure through endpoint marginal matching, transport-cost minimization, and cycle consistency.

Flow Matching and Coupling Design. Flow matching trains continuous normalizing flows by regressing velocity fields along prescribed probability paths [27]. The geometry of these paths depends critically on the coupling between source and target samples. Rectified flow and Reflow aim to learn straighter trajectories, but often require repeated simulation or retraining stages [4]. Minibatch OT and multisample couplings improve trajectory straightness by matching source and target samples within each batch [12, 13], but introduce batch-dependent approximations and additional OT computation. Recent semi-discrete and structured coupling methods further improve scalability or few-step generation by replacing batch OT with semi-discrete or hierarchical coupling mechanisms [22, 23, 24, 25]. In contrast, CyclOT does not require precomputed couplings, minibatch OT, or semi-discrete pairing.

Single-Step Generative Models. Accelerating generative sampling to one or a few steps is a major goal in diffusion and flow-based modeling. Consistency models and consistency trajectory models learn self-consistent maps along probability-flow trajectories [28, 29], while progressive distillation and shortcut models reduce the number of sampling steps by distilling or shortcutting iterative samplers [30, 31]. Recent mean-flow and flow-map-matching approaches further seek direct map-based generation [32, 33]. These approaches primarily focus on sampling acceleration, often through distillation, consistency constraints, or specialized schedules. Our goal is different: we seek to recover the quadratic OT map and its associated displacement interpolation from unpaired samples, thereby obtaining one-step generation as a consequence of OT map learning.

Unpaired Translation and Cycle Consistency. Cycle-consistent adversarial methods learn mappings between unpaired domains by combining adversarial distribution matching with inverse-consistency constraints [1]. These methods are highly influential in image translation, but they do not in general identify the quadratic OT map or the Wasserstein geodesic between distributions.

Our use of cycle consistency is instead tied to OT map recovery: the forward and backward maps define complementary displacement interpolants, while cycle consistency serves as a soft regularization aims to promote an approximately inverse relationship between the two maps and thereby imposes structural coherence on the induced transport paths.

1.3 Organization

Our paper is organized as follows. Section 2 is concerned with a review of the quadratic optimal transport formulation and the associated McCann displacement interpolation, which provide the foundation for our forward-backward transport framework. In Section 3, we propose our bidirectional map-based optimal transport framework, CyclOT, and introduce the corresponding objective function in the theoretical framework and its neural estimations that we use in the practical setting. In Section 4, we further prove that, under suitable uniform approximation assumptions on the generator and adversarial network classes, minimizers of the proposed objective converge to the quadratic optimal transport maps. We provide a local analysis of the proposed framework and further investigate the role of each term in the objective in Section 5. In particular, we study how cycle consistency restricts the feasible class of transport maps, how the action and cycle terms interact near inverse-compatible solutions, and how endpoint matching and cycle consistency suppress distinct forms of mode collapse. Section 6 is devoted to evaluating the CyclOT on five datasets: *Synthetic Swiss Roll*, *MNIST*, *CelebA*, *single-cell data*, and *chest X-ray images*. We compare CyclOT against five representative baselines: entropic OT (Sinkhorn) [11], Neural OT [16], ICNN-based OT [18], rectified flow [4], and mini-batch OT-guided flow matching [13]. Our experiments demonstrate that the proposed framework enables scalable and stable learning of optimal transport maps.

2 Problem Setup & Notation

2.1 Quadratic optimal transport

For a detailed introduction to optimal transport we recommend [5]. Let $\mathcal{P}_2(\mathbb{R}^d)$ denote the set of Borel probability measures on $\mathbb{R}^d$ with finite second moment. For $\mu, \nu \in \mathcal{P}_2(\mathbb{R}^d)$, let $\Pi(\mu, \nu)$ denote the set of their couplings. The quadratic Wasserstein distance is defined by

$$W_2^2(\mu, \nu) := \inf_{\pi \in \Pi(\mu, \nu)} \int_{\mathbb{R}^d \times \mathbb{R}^d} \|x - y\|^2 d\pi(x, y).$$

Throughout the paper, $\|\cdot\|$ denotes the Euclidean norm and Id denotes the identity map. For a measurable map $T : \mathbb{R}^d \rightarrow \mathbb{R}^d$, the pushforward $T_{\#}\mu$ is defined by

$$(T_{\#}\mu)(A) = \mu(T^{-1}(A))$$

for every Borel set $A \subset \mathbb{R}^d$. Equivalently, $T_{\#}\mu = \nu$ means that $T(X) \sim \nu$ whenever $X \sim \mu$.

Given $\mu_0, \mu_1 \in \mathcal{P}_2(\mathbb{R}^d)$, the Monge problem for the quadratic cost seeks a measurable map transporting μ_0 to μ_1 while minimizing

$$C_{\mu_0}(T) := \frac{1}{2} \int_{\mathbb{R}^d} \|T(x) - x\|^2 d\mu_0(x).$$

Thus, the quadratic Monge problem is

$$\inf_{\substack{T: \mathbb{R}^d \rightarrow \mathbb{R}^d \text{ measurable} \\ T_{\#}\mu_0 = \mu_1}} C_{\mu_0}(T). \tag{1}$$

Note that the factor $1/2$ is included in the Monge cost but not in the definition of W_2^2 . If μ_0 is absolutely continuous with respect to Lebesgue measure, Brenier's theorem gives a unique minimizer T^* , up to μ_0 -almost-everywhere equality, of the form

$$T^* = \nabla \phi \quad \mu_0\text{-a.e.},$$

where the potential function $\phi : \mathbb{R}^d \rightarrow (-\infty, +\infty]$ is convex [8, 5]. In particular,

$$T_{\#}^* \mu_0 = \mu_1, \quad C_{\mu_0}(T^*) = \frac{1}{2} W_2^2(\mu_0, \mu_1).$$

When μ_1 is also absolutely continuous, the reverse quadratic Monge problem admits a unique Brenier map S^* , up to μ_1 -almost-everywhere equality, satisfying

$$S_{\#}^* \mu_1 = \mu_0.$$

Suitable representatives of the two maps satisfy

$$S^* \circ T^* = \text{Id} \quad \mu_0\text{-a.e.}, \quad T^* \circ S^* = \text{Id} \quad \mu_1\text{-a.e.},$$

where Id denotes the identity map on the corresponding space.

Moreover,

$$C_{\mu_1}(S^*) = \frac{1}{2} W_2^2(\mu_0, \mu_1),$$

and hence

$$C_{\mu_0}(T^*) + C_{\mu_1}(S^*) = W_2^2(\mu_0, \mu_1). \quad (2)$$

We refer to (T^*, S^*) as the forward–reverse Brenier pair.

For the population recovery results, we work under the following standing setting. Let $\mathcal{X} \subset \mathbb{R}^d$ be a nonempty compact convex set with nonempty interior, and assume that

$$\mu_0, \mu_1 \in \mathcal{P}_{2,\text{ac}}(\mathcal{X}),$$

where

$$\mathcal{P}_{2,\text{ac}}(\mathcal{X}) := \left\{ \mu \in \mathcal{P}_2(\mathbb{R}^d) : \mu(\mathcal{X}) = 1, \mu \ll \mathcal{L}^d \right\}.$$

Here $\mathcal{L}^d$ denotes the d -dimensional Lebesgue measure.

2.2 Displacement interpolation

The forward Brenier map induces the McCann displacement interpolation [9]

$$F_t^*(x) := (1-t)x + tT^*(x), \quad \mu_t^* := (F_t^*)_{\#} \mu_0, \quad t \in [0, 1]. \quad (3)$$

Convexity of $\mathcal{X}$ ensures that $F_t^*(x) \in \mathcal{X}$ whenever $x, T^*(x) \in \mathcal{X}$. The curve $(\mu_t^*)_{t \in [0,1]}$ is the constant-speed quadratic Wasserstein geodesic from μ_0 to μ_1 satisfying

$$W_2(\mu_s^*, \mu_t^*) = |t-s| W_2(\mu_0, \mu_1), \quad s, t \in [0, 1].$$

The same interpolation can be represented using the reverse Brenier map. Define

$$G_t^*(y) := (1-t)S^*(y) + ty, \quad t \in [0, 1].$$

The almost-everywhere inverse relations imply

$$(G_t^\star)_\# \mu_1 = (F_t^\star)_\# \mu_0 = \mu_t^\star, \quad t \in [0, 1]. \quad (4)$$

Indeed, if $y = T^\star(x)$, then

$$G_t^\star(y) = (1-t)S^\star(T^\star(x)) + tT^\star(x) = (1-t)x + tT^\star(x) = F_t^\star(x)$$

for μ_0 -almost every x . Thus, although $G_t^\star$ is constructed from the reverse map, it parameterizes the same path in the time direction from μ_0 to μ_1 .

Each particle in (3) travels along a straight line with constant velocity $T^\star(x) - x$. Consequently,

$$\int_0^1 \mathbb{E}_{X \sim \mu_0} \left[\frac{1}{2} |\partial_t F_t^\star(X)|^2 \right] dt = \frac{1}{2} \mathbb{E}_{X \sim \mu_0} [|T^\star(X) - X|^2] = \frac{1}{2} W_2^2(\mu_0, \mu_1).$$

This identity motivates the quadratic-action term in the proposed objective, while (4) provides the ideal synchronization property that the learned forward-backward construction seeks to recover.

2.3 Bidirectional learning problem

We observe independent, unpaired samples from μ_0 and μ_1 . Our goal is to learn a forward map f and a backward map g that approximate $T^\star$ and $S^\star$, respectively. In particular, we seek endpoint marginal agreement,

$$f_\# \mu_0 \approx \mu_1, \quad g_\# \mu_1 \approx \mu_0,$$

together with approximate inverse consistency,

$$g \circ f \approx \text{Id} \quad \text{on } \mu_0, \quad f \circ g \approx \text{Id} \quad \text{on } \mu_1.$$

Marginal matching and inverse consistency alone do not identify the Brenier pair: many invertible maps can transport one distribution to the other. The quadratic transport cost supplies the additional variational criterion that selects the minimum-cost pair. CyclOT therefore combines endpoint matching, two-sided cycle consistency, and bidirectional quadratic action.

For a measurable vector-valued function h , we use the notation

$$\|h\|_{L^2(\mu)} := \left(\int_{\mathbb{R}^d} \|h(x)\|^2 d\mu(x) \right)^{1/2}.$$

Accordingly, population recovery of the Brenier pair will be measured by

$$\|f - T^\star\|_{L^2(\mu_0)} + \|g - S^\star\|_{L^2(\mu_1)}.$$

Section 3 turns the three structural requirements above into a neural population objective.

Additional background on flow matching, rectified flow, and independence-constrained Wasserstein barycenters is collected in Appendix A.

3 Proposed Framework

We introduce a bidirectional map-based framework for learning the quadratic optimal transport between μ_0 and μ_1 from independent endpoint samples. The framework consists of three neural approximation: a neural forward map, a neural reverse map, and neural endpoint-matching penalties, which together formulate the three training objectives: action, marginal matching, and cycle consistency. We first define the maps and their induced interpolants, then introduce the population objective and its practical minibatch implementation.

3.1 Forward-backward map framework

Let

$$\mathcal{F} = \{f_\theta : \theta \in \Theta\}, \quad \mathcal{G} = \{g_\phi : \phi \in \Phi\}$$

be admissible classes of maps from $\mathcal{X}$ to $\mathcal{X}$. We interpret f_θ as a forward map from μ_0 toward μ_1 , and g_ϕ as a reverse map from μ_1 toward μ_0 .

For $f \in \mathcal{F}$ and $g \in \mathcal{G}$, define the induced displacement interpolants

$$F_t^f(x) := (1-t)x + tf(x), \quad \mu_t^f := (F_t^f)_\# \mu_0,$$

and

$$G_t^g(y) := (1-t)g(y) + ty, \quad \nu_t^g := (G_t^g)_\# \mu_1, \quad t \in [0, 1].$$

Their endpoint laws are

$$\mu_0^f = \mu_0, \quad \mu_1^f = f_\# \mu_0,$$

and

$$\nu_0^g = g_\# \mu_1, \quad \nu_1^g = \mu_1.$$

Both paths are therefore oriented from μ_0 to μ_1 , although they are generated from opposite endpoints. Convexity of $\mathcal{X}$ ensures that the interpolated points remain in $\mathcal{X}$ whenever f and g are $\mathcal{X}$ -valued.

The two intermediate laws need not agree for arbitrary f and g . However, if

$$f_\# \mu_0 = \mu_1, \quad g \circ f = \text{Id} \quad \mu_0\text{-a.e.},$$

then

$$G_t^g(f(x)) = (1-t)g(f(x)) + tf(x) = (1-t)x + tf(x) = F_t^f(x)$$

for μ_0 -almost every x . Consequently,

$$\mu_t^f = \nu_t^g \quad \text{for every } t \in [0, 1].$$

Thus, the framework does not directly penalize discrepancies between intermediate marginals. Their synchronization follows from endpoint matching and forward-backward inverse consistency.

3.2 Neural population objective

We enforce the endpoint constraints using neural-approximation Jensen Shannon Divergence (neural JSD) objectives. Let $P, Q \in \mathcal{P}(\mathcal{X})$, where P is a generated distribution and Q is its target distribution. For a target-positive discriminator $D : \mathcal{X} \rightarrow (0, 1)$, define

$$\Phi_{P,Q}(D) := \log 2 + \frac{1}{2} \left[\int_{\mathcal{X}} \log(1-D) dP + \int_{\mathcal{X}} \log D dQ \right].$$

The unrestricted supremum of $\Phi_{P,Q}$ is the Jensen-Shannon divergence:

$$\text{JSD}(P, Q) = \sup_{\substack{D: \mathcal{X} \rightarrow (0,1) \\ D \text{ measurable}}} \Phi_{P,Q}(D).$$

Let $\mathcal{D}_m \subset C(\mathcal{X}; (0, 1))$ be a neural network class of capacity level m , and assume that $D \equiv 1/2 \in \mathcal{D}_m$. Define the neural JSD with capacity level m :

$$\widehat{\text{JSD}}_m(P, Q) := \sup_{D \in \mathcal{D}_m} \Phi_{P,Q}(D). \tag{5}$$

Note that, because $\Phi_{P,Q}(\frac{1}{2}) = 0$, we have

$$0 \leq \widehat{\text{JSD}}_m(P, Q) \leq \text{JSD}(P, Q).$$

The bidirectional neural endpoint penalty is

$$\widehat{\mathcal{B}}_m(f, g) := \widehat{\text{JSD}}_m(f_{\#}\mu_0, \mu_1) + \widehat{\text{JSD}}_m(g_{\#}\mu_1, \mu_0). \quad (6)$$

The two terms in (6) are implemented using separate forward and reverse neural JSD with the same capacity level m .

Furthermore, we define the bidirectional quadratic action by

$$\mathcal{A}(f, g) := \frac{1}{2} \int_{\mathcal{X}} \|f(x) - x\|^2 d\mu_0(x) + \frac{1}{2} \int_{\mathcal{X}} \|g(y) - y\|^2 d\mu_1(y), \quad (7)$$

and the two-sided cycle loss by

$$\mathcal{C}(f, g) := \int_{\mathcal{X}} \|g(f(x)) - x\|^2 d\mu_0(x) + \int_{\mathcal{X}} \|f(g(y)) - y\|^2 d\mu_1(y). \quad (8)$$

For an action coefficient $\lambda > 0$, the neural population objective is

$$\mathcal{E}_{\lambda, m}(f, g) := \lambda \mathcal{A}(f, g) + \widehat{\mathcal{B}}_m(f, g) + \mathcal{C}(f, g). \quad (9)$$

Here, the three terms have complementary roles:

- The endpoint term $\widehat{\mathcal{B}}_m$ promotes the forward and reverse marginal constraints via the neural JSD family with capacity level m

$$f_{\#}\mu_0 \approx \mu_1, \quad g_{\#}\mu_1 \approx \mu_0.$$

- The cycle term $\mathcal{C}$ explicitly promotes approximate inverse consistency in both directions for more stable and efficient learning by avoid mode collapses.
- The action term $\mathcal{A}$ favors maps with small quadratic displacement and supplies the optimal-transport selection criterion among approximately endpoint-feasible inverse pairs.

The action coefficient λ controls the action term relative to endpoint matching and cycle consistency. Small positive values allow the feasibility terms to dominate while retaining the quadratic-cost criterion.

3.3 Practical optimization algorithm

In practice, we choose f_{θ} and g_{ϕ} to be represented by neural networks. We use two target-positive discriminators: $D_{\omega_1}^1$ distinguishes samples from μ_1 from forward generated samples, while $D_{\omega_0}^0$ distinguishes samples from μ_0 from reverse generated samples.

At iteration k , draw independent minibatches

$$\{x_0^i\}_{i=1}^B \stackrel{\text{i.i.d.}}{\sim} \mu_0, \quad \{x_1^i\}_{i=1}^B \stackrel{\text{i.i.d.}}{\sim} \mu_1.$$

The empirical forward and reverse discriminator scores are

$$\begin{aligned}\widehat{\mathcal{V}}_B^1(\theta, \omega_1) &:= \frac{1}{2B} \sum_{i=1}^B [\log D_{\omega_1}^1(x_1^i) + \log(1 - D_{\omega_1}^1(f_\theta(x_0^i)))] , \\ \widehat{\mathcal{V}}_B^0(\phi, \omega_0) &:= \frac{1}{2B} \sum_{i=1}^B [\log D_{\omega_0}^0(x_0^i) + \log(1 - D_{\omega_0}^0(g_\phi(x_1^i)))] .\end{aligned}\tag{10}$$

The discriminator parameters are updated by maximizing these scores, or equivalently by minimizing $\widehat{\mathcal{L}}_{D^1} = -\widehat{\mathcal{V}}_B^1$ and $\widehat{\mathcal{L}}_{D^0} = -\widehat{\mathcal{V}}_B^0$.

For the map update, we use the non-saturating adversarial loss

$$\widehat{\mathcal{L}}_{\text{adv}}^{\text{NS}} := -\frac{1}{2B} \sum_{i=1}^B [\log D_{\omega_1}^1(f_\theta(x_0^i)) + \log D_{\omega_0}^0(g_\phi(x_1^i))].\tag{11}$$

This is a practical surrogate for the minimax neural-JSD map update and has the same intended marginal-matching equilibrium in the idealized expressive-discriminator setting.

The empirical action and cycle losses are

$$\widehat{\mathcal{A}}_B(\theta, \phi) := \frac{1}{2B} \sum_{i=1}^B [\|f_\theta(x_0^i) - x_0^i\|^2 + \|g_\phi(x_1^i) - x_1^i\|^2]\tag{12}$$

and

$$\widehat{\mathcal{C}}_B(\theta, \phi) := \frac{1}{B} \sum_{i=1}^B [\|g_\phi(f_\theta(x_0^i)) - x_0^i\|^2 + \|f_\theta(g_\phi(x_1^i)) - x_1^i\|^2].\tag{13}$$

The factor 1/2 appears in $\widehat{\mathcal{A}}_B$, consistently with (7), but not in $\widehat{\mathcal{C}}_B$, because (8) is the sum of the two directional cycle errors.

We use the action-weight schedule

$$\lambda_k = \lambda_{\text{init}} \beta^k, \quad \lambda_{\text{init}} > 0, \quad 0 < \beta < 1, \quad k = 0, \dots, K-1.\tag{14}$$

The map parameters are updated by minimizing

$$\widehat{\mathcal{L}}_{\text{map},k} = \widehat{\mathcal{L}}_{\text{adv}}^{\text{NS}} + \widehat{\mathcal{C}}_B + \lambda_k \widehat{\mathcal{A}}_B.\tag{15}$$

The complete procedure, which alternates between discriminator ascent and map descent, is summarized in Algorithm 1.

For direct consistency with the compact-domain analysis, the generator outputs should remain in $\mathcal{X}$. This can be enforced through a range-preserving output parameterization or, when convenient, by composing raw neural networks with the metric projection onto the compact convex set $\mathcal{X}$.

Theorem 4.4 concerns global minimizers of the fixed population minimax objective $\mathcal{E}_{\lambda,m}$. It does not establish convergence of the alternating stochastic algorithm, finite-sample discriminator training, the non-saturating surrogate (11), or the changing-weight schedule (14). These practical choices are motivated by the population objective but require a separate optimization analysis.

Algorithm 1 Forward-Backward Neural Optimal Transport

Require: Endpoint distributions μ_0, μ_1 ; generators f_θ, g_ϕ ; discriminators $D_{\omega_1}^1, D_{\omega_0}^0$; batch size B ; number of iterations K ; initial action weight $\lambda_{\text{init}} > 0$; annealing factor $\beta \in (0, 1)$.

1: **for** $k = 0, \dots, K - 1$ **do**

2: Draw independent minibatches

$$\{x_0^i\}_{i=1}^B \sim \mu_0, \quad \{x_1^i\}_{i=1}^B \sim \mu_1.$$

3: Compute $\hat{\mathcal{V}}_B^1$ and $\hat{\mathcal{V}}_B^0$ using (10).

4: Update ω_1, ω_0 by discriminator ascent:

$$\omega_1 \leftarrow \omega_1 + \eta_D \nabla_{\omega_1} \hat{\mathcal{V}}_B^1, \quad \omega_0 \leftarrow \omega_0 + \eta_D \nabla_{\omega_0} \hat{\mathcal{V}}_B^0.$$

5: Set

$$\lambda_k = \lambda_{\text{init}} \beta^k.$$

6: Compute $\hat{\mathcal{L}}_{\text{adv}}^{\text{NS}}, \hat{\mathcal{A}}_B$, and $\hat{\mathcal{C}}_B$.

7: Form

$$\hat{\mathcal{L}}_{\text{map},k} = \hat{\mathcal{L}}_{\text{adv}}^{\text{NS}} + \hat{\mathcal{C}}_B + \lambda_k \hat{\mathcal{A}}_B.$$

8: Update θ, ϕ by map descent:

$$\theta \leftarrow \theta - \eta_G \nabla_\theta \hat{\mathcal{L}}_{\text{map},k}, \quad \phi \leftarrow \phi - \eta_G \nabla_\phi \hat{\mathcal{L}}_{\text{map},k}.$$

9: **end for**

10: **return** f_θ, g_ϕ .

4 Global Recovery of the Quadratic Brenier Maps

Section 3 introduced the neural population objective

$$\mathcal{E}_{\lambda,m}(f, g) = \lambda \mathcal{A}(f, g) + \hat{\mathcal{B}}_m(f, g) + \mathcal{C}(f, g).$$

We now show that its global minimizers recover the forward and reverse quadratic Brenier maps when the discriminator level becomes sufficiently large and the action weight becomes sufficiently small but strictly positive.

The proof has three components. First, neural approximation provides a Brenier-approximating comparison sequence whose action converges to the optimal value and whose cycle loss vanishes. Second, an increasing discriminator family converts small neural JSD into uniform Wasserstein control of the generated marginals. Third, stability of the quadratic Monge problem converts approximate marginal feasibility and near-optimal action into L^2 -closeness to the Brenier maps.

4.1 Neural approximation ingredients

4.1.1 Generator approximation & range preservation

Classical universal-approximation results imply that suitable scalar neural families are dense in $C(K)$ under the uniform norm for every compact $K \subset \mathbb{R}^d$, cf. [34]. For example, Hornik [35] establishes uniform approximation on compact sets for continuous, bounded, nonconstant activations, while Leshno et al. [36] characterize universal approximation, under suitable regularity assumptions,

by nonpolynomiality of the activation. Applying these results coordinatewise gives local uniform approximation of continuous maps $K \rightarrow \mathbb{R}^d$.

Uniform approximation also implies L^p -approximation for every finite Borel measure μ on K and $1 \leq p < \infty$. Indeed, if $h_c \in C(K)$, then

$$\|N - h_c\|_{L^p(\mu)} \leq \mu(K)^{1/p} \|N - h_c\|_\infty.$$

Together with the fact that $C(K)$ is dense in $L^p(\mu)$ and triangle inequality, we have the neural family is dense in $L^p(\mu)$. The vector-valued conclusion follows coordinatewise.

We must additionally ensure that admissible generators take values in $\mathcal{X}$. Let $\mathcal{N}_f, \mathcal{N}_g \subset C(\mathbb{R}^d; \mathbb{R}^d)$ be locally uniformly dense raw neural classes, and define their range-preserving subclasses by

$$\mathcal{F} := \{N|_{\mathcal{X}} : N \in \mathcal{N}_f, N(\mathcal{X}) \subseteq \mathcal{X}\}, \quad \mathcal{G} := \{N|_{\mathcal{X}} : N \in \mathcal{N}_g, N(\mathcal{X}) \subseteq \mathcal{X}\}.$$

These subclasses remain uniformly dense in $C(\mathcal{X}; \mathcal{X})$. To see this, fix $x_0 \in \text{int}(\mathcal{X})$ and choose $\rho > 0$ such that $B(x_0, \rho) \subset \mathcal{X}$. For $h \in C(\mathcal{X}; \mathcal{X})$ and $0 < \delta < 1$, set

$$h_\delta := (1 - \delta)h + \delta x_0.$$

Convexity gives

$$B(h_\delta(x), \delta\rho) \subset \mathcal{X} \quad \text{for every } x \in \mathcal{X}.$$

Indeed, if $\|z - h_\delta(x)\| < \delta\rho$, then

$$z = (1 - \delta)h(x) + \delta \left(x_0 + \frac{z - h_\delta(x)}{\delta} \right) \in \mathcal{X}.$$

Choose a raw neural map N with

$$\|N - h_\delta\|_\infty < \frac{\delta\rho}{2}.$$

Then $N(\mathcal{X}) \subseteq \mathcal{X}$, while

$$\|N - h\|_\infty \leq \frac{\delta\rho}{2} + \delta \text{diam}(\mathcal{X}) \longrightarrow 0 \quad \text{as } \delta \downarrow 0.$$

Thus, $\mathcal{F}$ and $\mathcal{G}$ are uniformly dense in $C(\mathcal{X}; \mathcal{X})$. Moreover, convexity ensures that the straight-line interpolation between $x \in \mathcal{X}$ and $f(x) \in \mathcal{X}$ remains in $\mathcal{X}$.

4.1.2 Joint map & cycle approximation

To approximate an almost-everywhere inverse pair (T, S) , separate L^2 -approximations are not sufficient. Indeed, even if $T_j \rightarrow T$ in $L^2(\mu_0)$ and $S_j \rightarrow S$ in $L^2(\mu_1)$, it does not generally follow that $S_j \circ T_j \rightarrow S \circ T$ or $T_j \circ S_j \rightarrow T \circ S$ in the corresponding L^2 spaces. Consequently, the cycle loss need not converge to zero. The following lemma constructs a single approximating sequence that controls both map-approximation errors and the cycle loss simultaneously.

Lemma 4.1 (Joint approximation of an almost-everywhere inverse pair). *Let $\mathcal{X} \subset \mathbb{R}^d$ be a nonempty compact convex set, let $\mu_0, \mu_1 \in \mathcal{P}(\mathcal{X})$, and let $T, S : \mathcal{X} \rightarrow \mathcal{X}$ be measurable maps satisfying*

$$T_{\#}\mu_0 = \mu_1, \quad S \circ T = \text{Id} \quad \mu_0\text{-a.e.}, \quad T \circ S = \text{Id} \quad \mu_1\text{-a.e.}$$

Assume that $\mathcal{F}, \mathcal{G} \subset C(\mathcal{X}; \mathcal{X})$ are uniformly dense in $C(\mathcal{X}; \mathcal{X})$. Then there exists a sequence $(T_j, S_j) \in \mathcal{F} \times \mathcal{G}$ such that

$$\|T_j - T\|_{L^2(\mu_0)} + \|S_j - S\|_{L^2(\mu_1)} \longrightarrow 0$$

and

$$\mathcal{C}(T_j, S_j) \longrightarrow 0.$$

See Appendix B.1 for a proof of the lemma.

4.1.3 Discriminator uniform approximation

Universal approximation also yields discriminator classes that are dense in $C(\mathcal{X}; (0, 1))$. For continuous nonpolynomial hidden-layer activations, including ReLU, leaky ReLU, softplus, and ELU, this follows from Leshno et al. [36]. For continuous sigmoidal hidden-layer activations, such as the logistic sigmoid and hyperbolic tangent, see Cybenko [37] and Hornik [35].

Regardless of the hidden-layer activation, we use a logistic-sigmoid output. Thus every discriminator has the form

$$D = \sigma \circ \ell, \quad \sigma(u) := \frac{1}{1 + e^{-u}},$$

where $\ell : \mathcal{X} \rightarrow \mathbb{R}$ is a neural logit. In an implementation based on a binary-cross-entropy-with-logits loss, this sigmoid transformation may be applied implicitly for numerical stability.

For each $r \in \mathbb{N}$, let $\mathcal{L}_r$ be the union of the logit realization classes of finitely many neural architectures with compact parameter sets, and define

$$\mathcal{R}_r := \{\sigma \circ \ell : \ell \in \mathcal{L}_r\}.$$

Define the discriminator levels cumulatively by

$$\mathcal{D}_m := \{D \equiv 1/2\} \cup \bigcup_{r=1}^m \mathcal{R}_r.$$

Then

$$\mathcal{D}_1 \subseteq \mathcal{D}_2 \subseteq \dots, \quad D \equiv 1/2 \in \mathcal{D}_m,$$

and every $D \in \mathcal{D}_m$ takes values in $(0, 1)$.

Moreover, because $\mathcal{X}$ and the parameter sets at each fixed level are compact, and only finitely many architectures occur in $\mathcal{D}_m$, continuity of the sigmoid-output networks implies that there exists $c_m \in (0, 1/2)$ such that

$$c_m \leq D(x) \leq 1 - c_m \quad \text{for every } D \in \mathcal{D}_m \text{ and } x \in \mathcal{X}.$$

This fixed-level range bound is established explicitly in Lemma 4.2 below.

Suppose that the allowable architectures and parameter bounds increase with r , so that $\bigcup_{r \geq 1} \mathcal{L}_r$ is uniformly dense in $C(\mathcal{X})$. Let $H \in C(\mathcal{X}; (0, 1))$. Since $\mathcal{X}$ is compact, H is bounded away from zero and one, and hence

$$\text{logit}(H) := \log \frac{H}{1 - H} \in C(\mathcal{X}).$$

Given $\varepsilon > 0$, choose $\ell \in \bigcup_{r \geq 1} \mathcal{L}_r$ such that

$$\|\ell - \text{logit}(H)\|_\infty < 4\varepsilon.$$

Since the logistic sigmoid is $1/4$ -Lipschitz,

$$\|\sigma \circ \ell - H\|_\infty = \|\sigma \circ \ell - \sigma \circ \text{logit}(H)\|_\infty \leq \frac{1}{4} \|\ell - \text{logit}(H)\|_\infty < \varepsilon.$$

The network ℓ belongs to some finite level, so $\sigma \circ \ell \in \mathcal{D}_m$ for all sufficiently large m . Consequently,

$$\overline{\bigcup_{m \geq 1} \mathcal{D}_m} = C(\mathcal{X}; (0, 1))$$

in the relative uniform-norm topology.

4.1.4 Fixed-level discriminator compactness & continuity

At each fixed level, compact parameter sets and a uniform range bound yield compact discriminator and log-discriminator classes.

Lemma 4.2 (Fixed-level compactness of log-discriminator classes). *Fix $m \in \mathbb{N}$. Suppose that $\mathcal{D}_m$ is the realization class of finitely many neural architectures whose parameter sets are compact. Assume that the activations are continuous and that every realized discriminator is continuous and takes values in $(0, 1)$ (i.e. because it has a logistic-sigmoid output). Then, $\mathcal{D}_m$ is compact in*

$$(C(\mathcal{X}), \|\cdot\|_\infty),$$

and the families

$$\{\log D : D \in \mathcal{D}_m\}, \quad \{\log(1 - D) : D \in \mathcal{D}_m\}$$

are compact in the same space.

Proof. Index the finitely many architectures allowed at level m by $a = 1, \dots, A_m$. For architecture a , let $\Theta_{m,a}$ be its compact parameter set and define

$$F_{m,a} : \Theta_{m,a} \times \mathcal{X} \rightarrow (0, 1), \quad F_{m,a}(\theta, x) := D_{a,\theta}(x).$$

Because architecture a consists of finitely many affine layers composed with continuous activation functions, its realization

$$F_{m,a}(\theta, x) := D_{a,\theta}(x)$$

is jointly continuous in the parameters θ and the input x . Furthermore, since $\Theta_{m,a} \times \mathcal{X}$ is compact, $F_{m,a}$ is uniformly continuous. Consequently, the realization map

$$R_{m,a} : \Theta_{m,a} \rightarrow C(\mathcal{X}), \quad R_{m,a}(\theta) := D_{a,\theta},$$

is continuous with respect to the Euclidean norm on $\Theta_{m,a}$ and the uniform norm on $C(\mathcal{X})$. Indeed, if $\theta_n \rightarrow \theta$, uniform continuity gives

$$\|D_{a,\theta_n} - D_{a,\theta}\|_\infty = \sup_{x \in \mathcal{X}} |F_{m,a}(\theta_n, x) - F_{m,a}(\theta, x)| \rightarrow 0.$$

Now, since $\Theta_{m,a}$ is compact, its image $R_{m,a}(\Theta_{m,a})$ is compact in $C(\mathcal{X})$. Therefore,

$$\mathcal{D}_m = \bigcup_{a=1}^{A_m} R_{m,a}(\Theta_{m,a})$$

is compact as a finite union of compact sets. That proves the compactness of $\mathcal{D}_m$.

It remains to prove compactness of the logarithmic classes. For each architecture a , consider the continuous function

$$q_{m,a}(\theta, x) := \min\{F_{m,a}(\theta, x), 1 - F_{m,a}(\theta, x)\}.$$

Because every discriminator takes values strictly in $(0, 1)$, we have $q_{m,a}(\theta, x) > 0$ on the compact set $\Theta_{m,a} \times \mathcal{X}$. It then follows from the continuity of $q_{m,a}$ and compactness of $\Theta_{m,a} \times \mathcal{X}$ that

$$\eta_{m,a} := \min_{(\theta, x) \in \Theta_{m,a} \times \mathcal{X}} q_{m,a}(\theta, x) > 0.$$

Since only finitely many architectures are allowed, the number

$$c_m := \frac{1}{2} \min_{1 \leq a \leq A_m} \eta_{m,a}$$

is strictly positive and belongs to $(0, 1/2)$. It follows that $c_m \leq D(x) \leq 1 - c_m$ for every $D \in \mathcal{D}_m$ and $x \in \mathcal{X}$. On $[c_m, 1 - c_m]$, the scalar functions

$$u \mapsto \log u, \quad u \mapsto \log(1 - u)$$

are $1/c_m$ -Lipschitz. Thus, for $D, E \in \mathcal{D}_m$, we have

$$\|\log D - \log E\|_\infty \leq \frac{1}{c_m} \|D - E\|_\infty$$

and

$$\|\log(1 - D) - \log(1 - E)\|_\infty \leq \frac{1}{c_m} \|D - E\|_\infty.$$

Therefore, the maps

$$D \mapsto \log D, \quad D \mapsto \log(1 - D)$$

are continuous from $\mathcal{D}_m$ into $C(\mathcal{X})$. Their images of the compact set $\mathcal{D}_m$ are consequently compact. That completes the proof. $\square$

For $P, Q \in \mathcal{P}(\mathcal{X})$ and $D \in C(\mathcal{X}; (0, 1))$, define

$$\Phi_{P,Q}(D) := \log 2 + \frac{1}{2} \left[\int_{\mathcal{X}} \log D \, dP + \int_{\mathcal{X}} \log(1 - D) \, dQ \right].$$

Then, by definition,

$$\widehat{\text{JSD}}_m(P, Q) = \sup_{D \in \mathcal{D}_m} \Phi_{P,Q}(D).$$

We write $P_j \rightharpoonup P$ for weak convergence of probability measures. The following proposition shows that, at every fixed discriminator level m , the map

$$(P, Q) \mapsto \widehat{\text{JSD}}_m(P, Q)$$

is continuous on $\mathcal{P}(\mathcal{X}) \times \mathcal{P}(\mathcal{X})$ endowed with the product weak topology.

Proposition 4.3 (Fixed-level weak continuity of neural JSD). *Fix $m \in \mathbb{N}$, and suppose that*

$$\{\log D : D \in \mathcal{D}_m\}, \quad \{\log(1 - D) : D \in \mathcal{D}_m\}$$

are compact in $C(\mathcal{X})$. If $P_j, Q_j, P, Q \in \mathcal{P}(\mathcal{X})$ satisfy $P_j \rightharpoonup P$ and $Q_j \rightharpoonup Q$, then

$$\widehat{\text{JSD}}_m(P_j, Q_j) \longrightarrow \widehat{\text{JSD}}_m(P, Q).$$

If, in addition, $D \equiv 1/2 \in \mathcal{D}_m$, then

$$P_j \rightharpoonup Q \implies \widehat{\text{JSD}}_m(P_j, Q) \longrightarrow 0.$$

Proof. Let $\mathcal{H}_m^{(1)} := \{\log D : D \in \mathcal{D}_m\}$ and $\mathcal{H}_m^{(0)} := \{\log(1 - D) : D \in \mathcal{D}_m\}$. By Lemma 4.2, $\mathcal{H}_m^{(1)}$ is compact in $(C(\mathcal{X}), \|\cdot\|_\infty)$. Therefore, there exists a finite δ -net $\{h_1, \dots, h_N\}$ for $\mathcal{H}_m^{(1)}$. It follows that

$$\sup_{h \in \mathcal{H}_m^{(1)}} \left| \int h d(P_j - P) \right| \leq \max_{1 \leq i \leq N} \left| \int h_i d(P_j - P) \right| + 2\delta.$$

For each fixed $\delta > 0$, the functions $h_1, \dots, h_N$ are fixed. Weak convergence and finiteness of the net therefore imply

$$\max_{1 \leq i \leq N} \left| \int h_i d(P_j - P) \right| \longrightarrow 0.$$

Hence,

$$\limsup_{j \rightarrow \infty} \sup_{h \in \mathcal{H}_m^{(1)}} \left| \int h d(P_j - P) \right| \leq 2\delta.$$

Since $\delta > 0$ is arbitrary, letting $\delta \downarrow 0$ proves the desired convergence. The same argument applied to $\mathcal{H}_m^{(0)}$ gives

$$\sup_{D \in \mathcal{D}_m} \left| \int \log(1 - D) d(Q_j - Q) \right| \longrightarrow 0.$$

Consequently,

$$\sup_{D \in \mathcal{D}_m} |\Phi_{P_j, Q_j}(D) - \Phi_{P, Q}(D)| \longrightarrow 0.$$

Using

$$\left| \sup_D F_j(D) - \sup_D F(D) \right| \leq \sup_D |F_j(D) - F(D)|,$$

we obtain the first conclusion. For every $D \in \mathcal{D}_m$,

$$\Phi_{Q, Q}(D) = \log 2 + \frac{1}{2} \int_{\mathcal{X}} \log(D(1 - D)) dQ \leq 0,$$

because $D(1 - D) \leq 1/4$. The discriminator $D \equiv 1/2$ attains zero, and therefore

$$\widehat{\text{JSD}}_m(Q, Q) = 0.$$

The second conclusion follows from the first. □

We note here that the fixed-level qualification is essential. The proposition does not assert uniform convergence as $m \rightarrow \infty$.

4.2 Recovery of the forward-backward Brenier pair

Neural approximation properties used below. The preceding results yield the following two properties.

- (G) Since μ_0 and μ_1 are absolutely continuous, the forward and reverse quadratic Brenier maps are almost-everywhere inverses. The range-preserving generator construction and Lemma 4.1 therefore give a sequence

$$(T_j, S_j) \in \mathcal{F} \times \mathcal{G}$$

such that

$$\|T_j - T^*\|_{L^2(\mu_0)} + \|S_j - S^*\|_{L^2(\mu_1)} \longrightarrow 0, \quad \mathcal{C}(T_j, S_j) \longrightarrow 0.$$

- (D) The discriminator sequence satisfies $\mathcal{D}_1 \subseteq \mathcal{D}_2 \subseteq \dots$ and $D \equiv 1/2 \in \mathcal{D}_m$ for every fixed m , both log-discriminator families are compact in $C(\mathcal{X})$, and

$$\overline{\bigcup_{m \geq 1} \mathcal{D}_m} = C(\mathcal{X}; (0, 1))$$

in the relative uniform-norm topology.

Property (G) provides a Brenier map approximation whose action converges to W^2 and whose cycle loss vanishes. Property (D) provides fixed-level weak continuity and, through nestedness and density, uniform separation of distinct measures. No uniform approximation of true JSD and no simultaneous generator-discriminator diagonal are required.

For clarity, the neural population objective used below is

$$\begin{aligned} \widehat{\mathcal{B}}_m(f, g) &:= \widehat{\text{JSD}}_m(f_{\#}\mu_0, \mu_1) + \widehat{\text{JSD}}_m(g_{\#}\mu_1, \mu_0), \\ \mathcal{E}_{\lambda, m}(f, g) &:= \lambda \mathcal{A}(f, g) + \widehat{\mathcal{B}}_m(f, g) + \mathcal{C}(f, g). \end{aligned}$$

In the following theorem, the phrase “sufficiently expressive such that universal approximation holds” refers precisely to the concrete neural approximation constructions above and hence suffices properties (G) and (D).

Theorem 4.4 (OT Recovery under Neural JSD and Discriminator Annealing). *Let $\mathcal{X} \subset \mathbb{R}^d$ be a nonempty compact convex set with nonempty interior, and let $\mu_0, \mu_1 \in \mathcal{P}_{2,ac}(\mathcal{X})$. Let $T^*, S^* : \mathcal{X} \rightarrow \mathcal{X}$ be the almost-everywhere unique quadratic Brenier maps satisfying*

$$(T^*)_{\#}\mu_0 = \mu_1, \quad (S^*)_{\#}\mu_1 = \mu_0,$$

and set $W := W_2(\mu_0, \mu_1) > 0$. Assume that the generator classes $\mathcal{F}, \mathcal{G}$ and the discriminator class $\mathcal{D} := \bigcup_{m \geq 1} \mathcal{D}_m$ are sufficiently expressive such that property (G) and (D) are satisfied. For each $m \in \mathbb{N}$ and $\lambda > 0$, assume that the minimum is attained, and let

$$(f_{\lambda, m}, g_{\lambda, m}) \in \arg \min_{f \in \mathcal{F}, g \in \mathcal{G}} \mathcal{E}_{\lambda, m}(f, g).$$

Then, for every $\varepsilon > 0$, there exist $M_\varepsilon \in \mathbb{N}$ and $\alpha_\varepsilon > 0$ such that, whenever

$$m \geq M_\varepsilon \quad \text{and} \quad 0 < \lambda \leq \frac{\alpha_\varepsilon}{W^2},$$

one has

$$\|f_{\lambda, m} - T^*\|_{L^2(\mu_0)} + \|g_{\lambda, m} - S^*\|_{L^2(\mu_1)} < \varepsilon, \quad \mathcal{C}(f_{\lambda, m}, g_{\lambda, m}) \leq \lambda W^2.$$

Proof. We divide the proof into four steps.

Step 1: Brenier approximation benchmark. By property (G) which follows from Lemma 4.1, there exists $(T_j, S_j) \in \mathcal{F} \times \mathcal{G}$ such that

$$T_j \rightarrow T^\star \quad \text{in } L^2(\mu_0), \quad S_j \rightarrow S^\star \quad \text{in } L^2(\mu_1),$$

and $\mathcal{C}(T_j, S_j) \rightarrow 0$. The couplings $(T_j, T^\star)_{\#}\mu_0$ and $(S_j, S^\star)_{\#}\mu_1$ give

$$W_2((T_j)_{\#}\mu_0, \mu_1) \leq \|T_j - T^\star\|_{L^2(\mu_0)} \rightarrow 0$$

and

$$W_2((S_j)_{\#}\mu_1, \mu_0) \leq \|S_j - S^\star\|_{L^2(\mu_1)} \rightarrow 0.$$

Therefore, for every fixed m , Proposition 4.3 gives

$$\widehat{\mathcal{B}}_m(T_j, S_j) \rightarrow 0. \tag{16}$$

Here, m is held fixed while $j \rightarrow \infty$.

The same L^2 -convergence gives

$$\mathcal{A}(T_j, S_j) \rightarrow W^2.$$

Indeed,

$$\begin{aligned} & \left| \|T_j - \text{Id}\|_{L^2(\mu_0)}^2 - \|T^\star - \text{Id}\|_{L^2(\mu_0)}^2 \right| \\ & \leq (\|T_j - \text{Id}\|_{L^2(\mu_0)} + \|T^\star - \text{Id}\|_{L^2(\mu_0)}) \|T_j - T^\star\|_{L^2(\mu_0)} \rightarrow 0. \end{aligned}$$

and similarly for S_j . Moreover, $\mathcal{A}(T^\star, S^\star) = W^2$. Hence, we have

$$\mathcal{A}(T_j, S_j) \rightarrow \mathcal{A}(T^\star, S^\star) = W^2$$

by definition of $\mathcal{A}$.

Now, fix m and $\lambda > 0$. Global optimality of $(f_{\lambda,m}, g_{\lambda,m})$ gives

$$\mathcal{E}_{\lambda,m}(f_{\lambda,m}, g_{\lambda,m}) \leq \mathcal{E}_{\lambda,m}(T_j, S_j)$$

for every j . Letting $j \rightarrow \infty$ yields

$$\mathcal{E}_{\lambda,m}(f_{\lambda,m}, g_{\lambda,m}) \leq \lambda W^2. \tag{17}$$

Because $D \equiv 1/2 \in \mathcal{D}_m$, neural JSD is nonnegative. All three terms in $\mathcal{E}_{\lambda,m}$ are therefore nonnegative, and (17) implies

$$\mathcal{A}(f_{\lambda,m}, g_{\lambda,m}) \leq W^2, \tag{18}$$

$$\mathcal{C}(f_{\lambda,m}, g_{\lambda,m}) \leq \lambda W^2, \tag{19}$$

and

$$\widehat{\text{JSD}}_m((f_{\lambda,m})_{\#}\mu_0, \mu_1) \leq \lambda W^2, \tag{20}$$

$$\widehat{\text{JSD}}_m((g_{\lambda,m})_{\#}\mu_1, \mu_0) \leq \lambda W^2. \tag{21}$$

Step 2: Identification and uniform separation. For $P, Q \in \mathcal{P}(\mathcal{X})$, we first show that

$$\left[\widehat{\text{JSD}}_r(P, Q) = 0 \text{ for every } r \geq 1 \right] \implies P = Q. \quad (22)$$

Suppose that the quantities on the left vanish. Then

$$\Phi_{P,Q}(D) \leq 0 \quad \text{for every } D \in \bigcup_{r \geq 1} \mathcal{D}_r.$$

Fix $\varphi \in C(\mathcal{X})$. For all sufficiently small positive and negative t ,

$$D_t := \frac{1}{2} + t\varphi$$

belongs to $C(\mathcal{X}; (0, 1))$. By density, choose $D_{t,k} \in \bigcup_r \mathcal{D}_r$ with $D_{t,k} \rightarrow D_t$ uniformly. Since D_t is bounded away from zero and one,

$$\log D_{t,k} \rightarrow \log D_t, \quad \log(1 - D_{t,k}) \rightarrow \log(1 - D_t)$$

uniformly. Therefore

$$\Phi_{P,Q}(D_t) \leq 0.$$

At $t = 0$, $D_0 \equiv 1/2$ and $\Phi_{P,Q}(D_0) = 0$. Hence $t = 0$ is a local maximum of $t \mapsto \Phi_{P,Q}(D_t)$, and

$$0 = \frac{d}{dt} \Phi_{P,Q}(D_t) \Big|_{t=0} = \int_{\mathcal{X}} \varphi dP - \int_{\mathcal{X}} \varphi dQ.$$

Since this holds for every $\varphi \in C(\mathcal{X})$, the Riesz representation theorem gives $P = Q$, proving (22).

We next establish uniform separation. That is, for every $Q \in \mathcal{P}(\mathcal{X})$ and $a > 0$, there exist $M_{Q,a} \in \mathbb{N}$ and $\beta_{Q,a} > 0$ such that

$$m \geq M_{Q,a}, \quad \widehat{\text{JSD}}_m(P, Q) \leq \beta_{Q,a} \implies W_2(P, Q) < a \quad (23)$$

for every $P \in \mathcal{P}(\mathcal{X})$.

We prove by contradiction. Assume this is false, choosing the discriminator threshold $1/n$ and the minimum level n would give $m_n \geq n$ and $P_n \in \mathcal{P}(\mathcal{X})$ such that

$$W_2(P_n, Q) \geq a, \quad \widehat{\text{JSD}}_{m_n}(P_n, Q) \leq \frac{1}{n}.$$

Because $\mathcal{X}$ is compact, the space $\mathcal{P}(\mathcal{X})$ is compact under weak convergence. Hence, after passing to a subsequence, which we do not relabel, there exists $P \in \mathcal{P}(\mathcal{X})$ such that

$$P_n \rightharpoonup P.$$

On the compact space $\mathcal{X}$, weak convergence implies W_2 -convergence. Therefore,

$$W_2(P, Q) = \lim_{n \rightarrow \infty} W_2(P_n, Q) \geq a.$$

Now, fix $r \in \mathbb{N}$. Eventually $m_n \geq r$, so the set inclusion nestedness gives

$$0 \leq \widehat{\text{JSD}}_r(P_n, Q) \leq \widehat{\text{JSD}}_{m_n}(P_n, Q) \leq \frac{1}{n}.$$

Also, for fixed r , the map

$$P' \mapsto \widehat{\text{JSD}}_r(P', Q)$$

is lower semicontinuous under weak convergence, because it is a supremum of weakly continuous functions of P' . Hence,

$$0 \leq \widehat{\text{JSD}}_r(P, Q) \leq \liminf_{n \rightarrow \infty} \widehat{\text{JSD}}_r(P_n, Q) = 0.$$

Thus, the level- r value vanishes for every r . By (22), $P = Q$, which contradicts $W_2(P, Q) \geq a$. This proves (23).

Step 3: Stability of the quadratic Brenier map. Let $\mu \in \mathcal{P}_{2,ac}(\mathcal{X})$, $\nu \in \mathcal{P}(\mathcal{X})$, and T be the quadratic Brenier map from μ to ν . Write $W_{\mu,\nu} := W_2(\mu, \nu)$. We claim that, for every $\tau > 0$, there exist $\eta_\tau, \delta_\tau > 0$ such that every measurable $h : \mathcal{X} \rightarrow \mathcal{X}$ satisfying

$$W_2(h_{\#}\mu, \nu) \leq \eta_\tau$$

and

$$\int_{\mathcal{X}} \|h(x) - x\|^2 d\mu(x) \leq W_{\mu,\nu}^2 + \delta_\tau$$

also satisfies

$$\|h - T\|_{L^2(\mu)} < \tau.$$

Suppose otherwise. Then, for some $\tau_0 > 0$, there would exist measurable maps $h_n : \mathcal{X} \rightarrow \mathcal{X}$ such that

$$W_2((h_n)_{\#}\mu, \nu) \leq \frac{1}{n},$$

$$I_n := \int_{\mathcal{X}} \|h_n(x) - x\|^2 d\mu(x) \leq W_{\mu,\nu}^2 + \frac{1}{n},$$

but

$$\|h_n - T\|_{L^2(\mu)} \geq \tau_0.$$

The graph coupling induced by h_n gives

$$W_2(\mu, (h_n)_{\#}\mu) \leq I_n^{1/2}.$$

It therefore follows from the triangle inequality that

$$W_{\mu,\nu} \leq I_n^{1/2} + \frac{1}{n}.$$

Together with the upper bound on I_n , this implies

$$I_n \longrightarrow W_{\mu,\nu}^2.$$

Define

$$\gamma_n := (\text{Id}, h_n)_{\#}\mu.$$

After passing to a subsequence, $\gamma_n \rightharpoonup \gamma$ in $\mathcal{P}(\mathcal{X} \times \mathcal{X})$. Its first marginal is μ , and its second marginal is ν because $(h_n)_{\#}\mu \rightarrow \nu$ in W_2 . Continuity of the quadratic cost on the compact product gives

$$\int_{\mathcal{X} \times \mathcal{X}} \|x - y\|^2 d\gamma(x, y) = \lim_{n \rightarrow \infty} I_n = W_{\mu,\nu}^2.$$

Thus, γ is an optimal quadratic coupling between μ and ν . Since μ is absolutely continuous, the optimal quadratic coupling is unique and is induced by the Brenier map T , which is unique μ -almost everywhere. Therefore,

$$\gamma = (\text{Id}, T)_{\#}\mu.$$

Fix $s > 0$ and choose $T_s \in C(\mathcal{X}; \mathbb{R}^d)$ with $\|T_s - T\|_{L^2(\mu)} < s$. Since $(x, y) \mapsto \|y - T_s(x)\|^2$ is bounded and continuous, let $n \rightarrow \infty$, we have

$$\int_{\mathcal{X}} \|h_n(x) - T_s(x)\|^2 d\mu(x) \longrightarrow \int_{\mathcal{X}} \|T(x) - T_s(x)\|^2 d\mu(x) < s^2.$$

Moreover,

$$\|h_n - T\|_{L^2(\mu)}^2 \leq 2\|h_n - T_s\|_{L^2(\mu)}^2 + 2\|T_s - T\|_{L^2(\mu)}^2.$$

Consequently,

$$\limsup_{n \rightarrow \infty} \|h_n - T\|_{L^2(\mu)}^2 \leq 4s^2.$$

Letting $s \downarrow 0$ contradicts $\|h_n - T\|_{L^2(\mu)} \geq \tau_0$. This proves our stability claim.

Step 4: Forward and reverse recovery. Apply Step 3 with $\tau = \varepsilon/2$ to

$$(\mu, \nu, T) = (\mu_0, \mu_1, T^*)$$

and

$$(\mu, \nu, T) = (\mu_1, \mu_0, S^*).$$

Let the resulting constants be

$$\eta_\varepsilon^f, \delta_\varepsilon^f, \quad \eta_\varepsilon^g, \delta_\varepsilon^g,$$

and define

$$a_\varepsilon := \min \left\{ \eta_\varepsilon^f, \eta_\varepsilon^g, \frac{\delta_\varepsilon^f}{2W}, \frac{\delta_\varepsilon^g}{2W}, \frac{W}{2} \right\}.$$

Apply (23) with $(Q, a) = (\mu_1, a_\varepsilon)$ and $(Q, a) = (\mu_0, a_\varepsilon)$. Denote the resulting constants by

$$M_\varepsilon^f, \beta_\varepsilon^f, \quad M_\varepsilon^g, \beta_\varepsilon^g,$$

and set

$$M_\varepsilon := \max\{M_\varepsilon^f, M_\varepsilon^g\}, \quad \alpha_\varepsilon := \min\{\beta_\varepsilon^f, \beta_\varepsilon^g\}.$$

Suppose

$$m \geq M_\varepsilon, \quad 0 < \lambda \leq \frac{\alpha_\varepsilon}{W^2}.$$

The bounds (20) and (21), followed by uniform separation, give

$$W_2((f_{\lambda,m})_\# \mu_0, \mu_1) < a_\varepsilon, \quad W_2((g_{\lambda,m})_\# \mu_1, \mu_0) < a_\varepsilon.$$

Write

$$I_f := \int_{\mathcal{X}} \|f_{\lambda,m}(x) - x\|^2 d\mu_0(x), \quad I_g := \int_{\mathcal{X}} \|g_{\lambda,m}(y) - y\|^2 d\mu_1(y).$$

By the action bound (18) established in Step 1 and the definition of the bidirectional action,

$$\frac{1}{2}(I_f + I_g) = \mathcal{A}(f_{\lambda,m}, g_{\lambda,m}) \leq W^2.$$

Therefore,

$$I_f + I_g \leq 2W^2.$$

Now, let $\nu_g := (g_{\lambda,m})_\# \mu_1$. The graph coupling $(\text{Id}, g_{\lambda,m})_\# \mu_1$ gives

$$W_2(\mu_1, \nu_g) \leq I_g^{1/2}.$$

Moreover, the reverse triangle inequality for W_2 gives

$$W_2(\mu_1, \nu_g) \geq W_2(\mu_1, \mu_0) - W_2(\nu_g, \mu_0).$$

Consequently,

$$I_g^{1/2} \geq W - W_2((g_{\lambda,m})_\# \mu_1, \mu_0) > W - a_\varepsilon.$$

Therefore,

$$I_f \leq 2W^2 - I_g < 2W^2 - (W - a_\varepsilon)^2 = W^2 + 2Wa_\varepsilon - a_\varepsilon^2 \leq W^2 + \delta_\varepsilon^f.$$

The symmetric argument gives

$$I_g < W^2 + \delta_\varepsilon^g.$$

Since $a_\varepsilon \leq \eta_\varepsilon^f$ and $a_\varepsilon \leq \eta_\varepsilon^g$, it follows from Step 3 that:

$$\|f_{\lambda,m} - T^*\|_{L^2(\mu_0)} < \frac{\varepsilon}{2}, \quad \|g_{\lambda,m} - S^*\|_{L^2(\mu_1)} < \frac{\varepsilon}{2}.$$

Adding these inequalities proves the map-recovery conclusion, while (19) gives

$$\mathcal{C}(f_{\lambda,m}, g_{\lambda,m}) \leq \lambda W^2.$$

□

4.3 Discriminator refinement & action-weight annealing

The transport weight λ balances the transportation cost against marginal matching and cycle consistency. The theoretical estimate gives

$$\widehat{\mathcal{B}}_m(f_{\lambda,m}, g_{\lambda,m}) + \mathcal{C}(f_{\lambda,m}, g_{\lambda,m}) \leq \lambda W^2, \quad \mathcal{A}(f_{\lambda,m}, g_{\lambda,m}) \leq W^2,$$

with increasing discriminator class indexed by m . Thus, for consistent implementation, annealing λ toward zero is essential; together with a sufficiently expressive discriminator level, it forces marginal mismatch to vanish to mimic the increasing discriminator class. At the same time, λ must remain strictly positive. Even when its coefficient is small, the action term retains the optimal-transport selection mechanism among approximately feasible transports. Finally, at the end of annealing, the matching has already been forged during the previous training dynamics when the actions scale is still large. Therefore, although the final training epochs should have nearly zero action terms and focuses on distribution matching, the perturbation on the already forged matching remains relatively small.

The proof also relates the annealing scale to the required accuracy of marginal enforcement. For a prescribed tolerance $a > 0$, the separation property provides discriminator levels and thresholds

$$M_{\mu_1,a}, \beta_{\mu_1,a}, \quad M_{\mu_0,a}, \beta_{\mu_0,a}.$$

Consequently, if

$$m \geq \max\{M_{\mu_1,a}, M_{\mu_0,a}\}, \quad \lambda W^2 \leq \min\{\beta_{\mu_1,a}, \beta_{\mu_0,a}\},$$

then both generated marginals lie within W_2 -distance a of their targets. Hence λW^2 acts as a feasibility-error budget, while the quantities $\beta_{\mu_i,a}$ describe the resolution of the discriminator family at the desired marginal scale. The theorem uses this relation with $a = a_\varepsilon$. This calibration is qualitative: an explicit annealing rate would require quantitative bounds on the discriminator separation modulus.

This suggests a continuation interpretation. A larger initial value of λ imposes a stronger low-displacement geometric bias. Gradually reducing λ then transfers emphasis to marginal matching and cycle consistency, while a residual positive action weight preserves the optimal-transport selection principle. The theorem establishes this small-positive- λ behavior for global minimizers; it does not assert that beginning with a large λ , or any particular annealing schedule, is necessary for optimization, nor does it analyze local minima.

5 Structural Roles of Endpoint Matching, Cycle Consistency, & Action

Section 4 establishes global recovery of the forward–reverse Brenier pair. We now examine the structural mechanisms behind this result. The analysis addresses three questions: how cycle consistency

reduces the class of admissible transport maps, how the proposed penalties control different forms of mode collapse, and how the local competition between quadratic action and cycle consistency disfavors suboptimal inverse-compatible solutions.

The discussion has three parts. First, endpoint matching alone leaves a large class of admissible forward–backward map pairs. Cycle consistency narrows this class by requiring the two maps to be approximately inverse to one another. This reduction is substantial, although it does not by itself identify the Brenier pair, because many inverse-compatible maps can transport the prescribed endpoint marginals.

Second, we distinguish two forms of mode collapse. Distribution-level collapse occurs when a generated endpoint distribution fails to cover a region carrying nonnegligible target mass. After sufficient discriminator refinement, the neural-JSD term controls this failure at any prescribed positive spatial scale. Map-level collapse occurs when many distinct inputs are sent to the same output. The cycle loss controls this second failure by bounding the conditional dispersion of inputs associated with a common generated output, as well as the corresponding dispersion of the induced trajectories.

Third, we study the local competition between quadratic action and cycle consistency near an inverse-compatible pair. Along a smooth endpoint-feasible perturbation that decreases action, the action improvement is generally first order in the perturbation size, whereas the cycle error created by incompatibility of the forward and backward perturbations is second order. Thus, the action term favors motion toward a lower-cost transport, while the cycle term resists departures from inverse compatibility. In particular, a suboptimal inverse-compatible pair admitting such an action-decreasing feasible perturbation cannot be a local minimizer of the population objective. At the Brenier pair, by contrast, no endpoint-feasible perturbation can further reduce the quadratic action. This provides a local structural interpretation of how the two terms jointly favor the Brenier pair; it is not, by itself, a convergence guarantee for gradient-based training.

Throughout this section, let $\mathcal{X} \subset \mathbb{R}^d$ be compact and convex, let $\mu_0, \mu_1 \in \mathcal{P}_{2,ac}(\mathcal{X})$, and let (T^*, S^*) denote the forward–reverse quadratic Brenier pair. We use the population objective introduced in Section 3:

$$\mathcal{E}_{\lambda,m}(f, g) = \lambda \mathcal{A}(f, g) + \widehat{\mathcal{B}}_m(f, g) + \mathcal{C}(f, g),$$

where

$$\widehat{\mathcal{B}}_m(f, g) = \widehat{\text{JSD}}_m(f_{\#}\mu_0, \mu_1) + \widehat{\text{JSD}}_m(g_{\#}\mu_1, \mu_0).$$

The three terms therefore play complementary roles:

- Cycle consistency narrows the admissible map class by promoting approximate forward–backward invertibility and controlling many-to-one map collapse.
- Neural-JSD endpoint matching controls discrepancies between the generated and target endpoint distributions after sufficient discriminator refinement.
- Quadratic action favors lower-cost transports and, in competition with the cycle penalty, provides the minimum-action selection mechanism leading to the forward–reverse Brenier pair.

At a fixed discriminator level m , vanishing neural JSD need not imply exact equality of the corresponding measures. Consequently, exact endpoint matching is used below only to describe the ideal limiting feasible class. Its connection to the neural objective is provided by the identification and uniform Wasserstein-separation properties established in Section 4.

5.1 Feasible-class reduction via cycle consistency

Define the exact endpoint-matching class

$$\mathcal{M}_{\text{match}} := \{(f, g) : f_{\#}\mu_0 = \mu_1, \quad g_{\#}\mu_1 = \mu_0\}.$$

Under the discriminator assumptions of Section 4, this limiting feasible class can also be characterized through the full discriminator sequence:

$$(f, g) \in \mathcal{M}_{\text{match}} \iff \widehat{\mathcal{B}}_r(f, g) = 0 \quad \text{for every } r \geq 1.$$

In general, however, vanishing neural JSD at one fixed discriminator level does not imply exact endpoint matching.

Endpoint matching alone leaves a large class of admissible map pairs. More strongly, even exact cycle consistency does not uniquely determine the Brenier pair. To see this, let $R : \mathcal{X} \rightarrow \mathcal{X}$ be a measurable μ_0 -preserving bijection with a measurable inverse, where these properties are understood up to μ_0 -null sets, and define

$$f_R := T^* \circ R, \quad g_R := R^{-1} \circ S^*.$$

Then

$$(f_R)_{\#}\mu_0 = \mu_1, \quad (g_R)_{\#}\mu_1 = \mu_0.$$

Moreover, because T^* and S^* are almost-everywhere inverses,

$$g_R \circ f_R = \text{Id} \quad \mu_0\text{-a.e.}, \quad f_R \circ g_R = \text{Id} \quad \mu_1\text{-a.e.}$$

Thus, endpoint matching together with exact cycle consistency still leaves a potentially large class of inverse-compatible transport pairs.

We next quantify the restriction imposed by a small cycle loss. For $\varepsilon > 0$, define the two-sided reconstruction-failure probability

$$p_\varepsilon(f, g) := \mathbb{P}_{X \sim \mu_0}(\|g(f(X)) - X\| > \varepsilon) + \mathbb{P}_{Y \sim \mu_1}(\|f(g(Y)) - Y\| > \varepsilon).$$

Proposition 5.1 (Approximate invertibility at resolution ε). *For every $\varepsilon > 0$,*

$$p_\varepsilon(f, g) \leq \min \left\{ 2, \frac{\mathcal{C}(f, g)}{\varepsilon^2} \right\}.$$

Consequently, if $\mathcal{C}(f, g) \leq \delta$, then

$$p_\varepsilon(f, g) \leq \min \left\{ 2, \frac{\delta}{\varepsilon^2} \right\}.$$

See Appendix B.2 for the proof of this statement. Proposition 5.1 shows that cycle consistency narrows the endpoint-matched candidate class by requiring the forward and backward maps to reconstruct one another with high probability at every fixed positive resolution. Nevertheless, the construction (f_R, g_R) shows that even exact cycle consistency leaves multiple inverse-compatible candidates. Cycle consistency is therefore a candidate-reduction mechanism rather than, by itself, an optimal-transport selection principle.

5.2 Endpoint mode coverage and suppression of map-level collapse

Mode collapse can refer to two distinct failures. First, the generated endpoint distribution may fail to cover a region carrying nonnegligible target mass. Second, the learned map may send widely dispersed inputs to the same output, making the input difficult to recover from the generated sample. Neural-JSD endpoint matching and cycle consistency control these two failures in complementary ways.

We first formulate the endpoint guarantee using the neural JSD appearing in the proposed objective. For a Borel set $A \subseteq \mathcal{X}$ and $r > 0$, define its r -neighborhood by

$$A^{(r)} := \{z \in \mathcal{X} : \text{dist}(z, A) < r\}.$$

Proposition 5.2 (Neural-JSD mode coverage at a prescribed spatial scale). *Assume the discriminator conditions of Section 4. The, for every $r > 0$ and $\delta \in (0, 1)$, there exist*

$$M_{r,\delta} \in \mathbb{N}, \quad \beta_{r,\delta} > 0,$$

such that, whenever

$$m \geq M_{r,\delta}, \quad \widehat{\mathcal{B}}_m(f, g) \leq \beta_{r,\delta},$$

one has, for every Borel set $A \subseteq \mathcal{X}$,

$$(f_{\#}\mu_0)(A^{(r)}) \geq \mu_1(A) - \delta$$

and

$$(g_{\#}\mu_1)(A^{(r)}) \geq \mu_0(A) - \delta.$$

In particular, if

$$(f_{\#}\mu_0)(A^{(r)}) = 0,$$

then $\mu_1(A) \leq \delta$, and the analogous conclusion holds in the reverse direction.

Proof. Set

$$a := r\sqrt{\delta}.$$

Apply the neural-JSD separation property from Section 4 to the target measures μ_1 and μ_0 . Taking the larger of the resulting discriminator levels and the smaller of the resulting neural-JSD thresholds gives $M_{r,\delta}$ and $\beta_{r,\delta}$ such that

$$m \geq M_{r,\delta}, \quad \widehat{\mathcal{B}}_m(f, g) \leq \beta_{r,\delta}$$

imply

$$W_2(f_{\#}\mu_0, \mu_1) < a$$

and

$$W_2(g_{\#}\mu_1, \mu_0) < a.$$

Let $P := f_{\#}\mu_0$, and let $\pi \in \Pi(P, \mu_1)$ be an optimal quadratic coupling. If $y \in A$ and $z \notin A^{(r)}$, then $\|z - y\| \geq r$. Therefore,

$$\begin{aligned} \mu_1(A) &\leq P(A^{(r)}) + \pi(\{(z, y) : \|z - y\| \geq r\}) \\ &\leq P(A^{(r)}) + \frac{1}{r^2} \int_{\mathcal{X} \times \mathcal{X}} \|z - y\|^2 d\pi(z, y) \\ &= P(A^{(r)}) + \frac{W_2^2(P, \mu_1)}{r^2} \\ &< P(A^{(r)}) + \delta. \end{aligned}$$

This proves

$$(f_{\#}\mu_0)(A^{(r)}) \geq \mu_1(A) - \delta.$$

The reverse inequality follows by applying the same argument to $g_{\#}\mu_1$ and μ_0 . $\square$

The enlargement from A to $A^{(r)}$ is essential. Wasserstein closeness does not control the probabilities of arbitrary sets without additional assumptions on their boundaries. Likewise, a small neural JSD at one fixed discriminator level does not imply the total-variation estimate available for unrestricted population JSD. Proposition 5.2 instead gives the appropriate neural-JSD guarantee: after sufficient discriminator refinement, target mass cannot be entirely absent from a prescribed positive neighborhood.

For the global minimizers considered in Theorem 4.4, the benchmark inequality gives

$$\widehat{\mathcal{B}}_m(f_{\lambda,m}, g_{\lambda,m}) \leq \lambda W_2^2(\mu_0, \mu_1).$$

Consequently, the conclusion of Proposition 5.2 holds whenever

$$m \geq M_{r,\delta}, \quad 0 < \lambda \leq \frac{\beta_{r,\delta}}{W_2^2(\mu_0, \mu_1)}.$$

This relates the spatial resolution r , the tolerated missing mass δ , the discriminator level m , and the action weight λ .

We next quantify the many-to-one map collapse. Let $X_0 \sim \mu_0$ and $X_1 \sim \mu_1$, and define

$$\mathcal{V}_0(f) := \mathbb{E}[\text{Var}(X_0 \mid f(X_0))], \quad \mathcal{V}_1(g) := \mathbb{E}[\text{Var}(X_1 \mid g(X_1))],$$

where

$$\text{Var}(X \mid Z) := \mathbb{E}[\|X - \mathbb{E}[X \mid Z]\|^2 \mid Z].$$

These quantities measure the average dispersion of inputs associated with a common generated output.

Define the directional cycle losses

$$\mathcal{C}_0(f, g) := \mathbb{E}\|g(f(X_0)) - X_0\|^2, \quad \mathcal{C}_1(f, g) := \mathbb{E}\|f(g(X_1)) - X_1\|^2.$$

Proposition 5.3 (Cycle consistency controls many-to-one collapse). *For any measurable maps f, g for which the displayed quantities are finite,*

$$\mathcal{V}_0(f) \leq \mathcal{C}_0(f, g), \quad \mathcal{V}_1(g) \leq \mathcal{C}_1(f, g).$$

Consequently,

$$\mathcal{V}_0(f) + \mathcal{V}_1(g) \leq \mathcal{C}(f, g).$$

Proof. The random variable $g(f(X_0))$ is measurable with respect to $\sigma(f(X_0))$. Since $\mathbb{E}[X_0 \mid f(X_0)]$ is the best $\sigma(f(X_0))$ -measurable L^2 -predictor of X_0 ,

$$\mathbb{E}\|X_0 - \mathbb{E}[X_0 \mid f(X_0)]\|^2 \leq \mathbb{E}\|X_0 - g(f(X_0))\|^2.$$

This proves the source-side inequality. Interchanging (X_0, f, g) with (X_1, g, f) proves the target-side inequality. Adding the two bounds gives the final claim. $\square$

The same estimate controls conditional trajectory dispersion. Define

$$X_t^F = (1-t)X_0 + tf(X_0), \quad X_t^B = (1-t)g(X_1) + tX_1.$$

Conditionally on $f(X_0)$, the term $tf(X_0)$ is fixed, and conditionally on $g(X_1)$, the term $(1-t)g(X_1)$ is fixed. Consequently,

$$\mathbb{E}[\text{Var}(X_t^F \mid f(X_0))] = (1-t)^2 \mathcal{V}_0(f) \leq (1-t)^2 \mathcal{C}_0(f, g)$$

and

$$\mathbb{E}[\text{Var}(X_t^B \mid g(X_1))] = t^2 \mathcal{V}_1(g) \leq t^2 \mathcal{C}_1(f, g).$$

Propositions 5.2 and 5.3 therefore describe complementary controls. After sufficient discriminator refinement, neural JSD prevents target mass from being absent at a prescribed positive spatial scale, whereas cycle consistency controls many-to-one preimage collapse and the resulting conditional dispersion of learned trajectories. Consistent with this analysis, the cycle ablation results in Appendix D.4 show that removing cycle consistency leads to mode collapse. These are population-level guarantees and do not, by themselves, eliminate errors caused by finite samples or incomplete optimization.

5.3 Local interaction among action, neural JSD, & cycle consistency

We next examine the local interaction among the three terms in the population objective defined in equation 9. Unlike an exactly constrained formulation, CyclOT does not project generator updates onto the marginal-matching or inverse-compatible classes. We therefore analyze a general perturbation of the neural-generator parameters.

Fix a discriminator level m , and write the generators as $f = f_\theta$ and $g = g_\phi$. An admissible smooth perturbation is a pair of one-sided C^2 parameter curves

$$\alpha \mapsto \theta_\alpha, \quad \alpha \mapsto \phi_\alpha, \quad \alpha \geq 0,$$

that remain in the admissible parameter sets and satisfy

$$\theta_0 = \theta, \quad \phi_0 = \phi.$$

Any range restrictions imposed on the generators are also required to hold along these curves.

For this subsection only, assume that the generator realization maps are jointly C^2 in their inputs and parameters on a neighborhood of the relevant compact sets. This holds, for example, for finite neural networks with C^2 activations and smooth output maps along bounded parameter curves. Writing

$$f_\alpha := f_{\theta_\alpha}, \quad g_\alpha := g_{\phi_\alpha},$$

we then have

$$f_\alpha = f - \alpha U_f + O(\alpha^2), \quad g_\alpha = g - \alpha U_g + O(\alpha^2),$$

with uniform C^1 remainders, where

$$U_f := - \left. \frac{d}{d\alpha} f_{\theta_\alpha} \right|_{\alpha=0+}, \quad U_g := - \left. \frac{d}{d\alpha} g_{\phi_\alpha} \right|_{\alpha=0+}.$$

Define

$$a(f, g; U_f, U_g) := \int_{\mathcal{X}} \langle f(x) - x, U_f(x) \rangle d\mu_0(x) + \int_{\mathcal{X}} \langle g(y) - y, U_g(y) \rangle d\mu_1(y).$$

The quadratic action satisfies

$$\mathcal{A}(f_\alpha, g_\alpha) = \mathcal{A}(f, g) - \alpha a(f, g; U_f, U_g) + O(\alpha^2).$$

Thus $a > 0$ means that the perturbation decreases the action to first order.

Because the perturbation may change the endpoint marginals, the neural-JSD term also contributes. Suppose that the fixed-level discriminator class has a compact parameterization, its discriminators are continuously differentiable on the relevant compact set, and their values are uniformly bounded away from 0 and 1. Danskin's theorem then gives the one-sided expansion

$$\widehat{\mathcal{B}}_m(f_\alpha, g_\alpha) = \widehat{\mathcal{B}}_m(f, g) + \alpha b_m(f, g; U_f, U_g) + o(\alpha),$$

where

$$b_m(f, g; U_f, U_g) := \left. \frac{d}{d\alpha} \widehat{\mathcal{B}}_m(f_\alpha, g_\alpha) \right|_{\alpha=0+}.$$

If the maximizing discriminator is not unique, this derivative is obtained by maximizing the corresponding directional derivative over the active maximizers.

We next expand the cycle loss without assuming inverse compatibility. Define the current residuals

$$R_0(x) := g(f(x)) - x, \quad R_1(y) := f(g(y)) - y,$$

and their first-order changes

$$W_0(x) := Dg(f(x))U_f(x) + U_g(f(x)),$$

$$W_1(y) := Df(g(y))U_g(y) + U_f(g(y)).$$

Taylor expansion gives

$$g_\alpha(f_\alpha(x)) - x = R_0(x) - \alpha W_0(x) + O(\alpha^2)$$

and

$$f_\alpha(g_\alpha(y)) - y = R_1(y) - \alpha W_1(y) + O(\alpha^2).$$

Consequently,

$$\mathcal{C}(f_\alpha, g_\alpha) = \mathcal{C}(f, g) - 2\alpha \Gamma(f, g; U_f, U_g) + O(\alpha^2),$$

where

$$\Gamma(f, g; U_f, U_g) := \langle R_0, W_0 \rangle_{L^2(\mu_0)} + \langle R_1, W_1 \rangle_{L^2(\mu_1)}.$$

Thus $\Gamma > 0$ means that the perturbation decreases cycle loss to first order, whereas $\Gamma < 0$ means that cycle consistency opposes the perturbation.

For

$$E_{\text{mis}}(f, g; U_f, U_g) := \|W_0\|_{L^2(\mu_0)}^2 + \|W_1\|_{L^2(\mu_1)}^2,$$

Cauchy–Schwarz gives

$$|\Gamma(f, g; U_f, U_g)| \leq \sqrt{\mathcal{C}(f, g) E_{\text{mis}}(f, g; U_f, U_g)}.$$

Hence, along any fixed perturbation, the first-order cycle contribution is small near a low-cycle-loss pair. In the special case of exact inverse compatibility, $R_0 = R_1 = 0$, so $\Gamma = 0$ and

$$\mathcal{C}(f_\alpha, g_\alpha) = \alpha^2 E_{\text{mis}}(f, g; U_f, U_g) + O(\alpha^3).$$

Exact inverse compatibility is therefore not assumed; it merely recovers the previous quadratic special case.

Combining the three expansions for

$$\mathcal{E}_{\lambda,m} = \lambda\mathcal{A} + \widehat{\mathcal{B}}_m + \mathcal{C}$$

yields

$$\begin{aligned} \mathcal{E}_{\lambda,m}(f_\alpha, g_\alpha) - \mathcal{E}_{\lambda,m}(f, g) &= \alpha [b_m(f, g; U_f, U_g) - \lambda a(f, g; U_f, U_g) - 2\Gamma(f, g; U_f, U_g)] \\ &\quad + o(\alpha). \end{aligned}$$

Therefore, if

$$\lambda a(f, g; U_f, U_g) + 2\Gamma(f, g; U_f, U_g) > b_m(f, g; U_f, U_g),$$

then the complete neural-JSD population objective strictly decreases for all sufficiently small $\alpha > 0$. Any pair admitting such a direction cannot be a local minimizer or a stationary point with respect to all admissible directions.

This condition separates the three local forces. The term λa favors lower quadratic transport cost, b_m measures the change in endpoint matching, and Γ measures whether the same update improves or worsens cycle consistency. When $\Gamma < 0$, action and cycle consistency compete; when $\Gamma > 0$, they cooperate. Decreasing λ correspondingly shifts the local balance from transport cost toward endpoint matching and cycle consistency.

This is a local population-level calculation with optimized discriminators. It does not imply that every nonoptimal pair admits a decreasing direction, nor does it establish convergence of finite-sample alternating or stochastic training. Global recovery of the Brenier pair remains the content of Theorem 4.4.

6 Experiments

We evaluate the proposed method, CyclOT, on five different datasets, spanning synthetic geometry, handwritten images, natural images, and scientific data: *Swiss roll*, *MNIST*, *CelebA*, *single-cell data*, and *chest X-ray images*. The experiments are designed to test three aspects of the method that are central to our claims: (i) *marginal matching*, namely whether the learned pushforward matches the target distribution; (ii) *transport structure*, namely whether the learned map remains short, regular, and approximately invertible; and (iii) *intermediate geometry*, namely whether the induced interpolation path is more coherent than that of standard baselines.

6.1 Experimental protocol & interpretation of metrics

We evaluate the three complementary roles represented in the population objective

$$\mathcal{E}_{\lambda,m} = \lambda\mathcal{A} + \widehat{\mathcal{B}}_m + \mathcal{C} :$$

endpoint marginal matching, quadratic transport cost, and approximate inverse consistency. Because the true Brenier maps are unavailable in the real-data experiments and there exists a trade-off between the transportation loss and the marginal distribution matching loss in practice, no single reported metric measures OT-map recovery performance. We therefore use several complementary diagnostics and interpret them jointly.

Comparison methods. We compare CyclOT with entropic OT based on Sinkhorn iterations [11], Neural OT [16], ICNN OT [18], rectified flow with reflow [4], and mini-batch OT-guided conditional flow matching (OT-CFM) [13]. Sinkhorn is treated as a finite-sample, transductive coupling reference that does not define an out-of-sample map. The remaining methods learn global neural maps or ODE transports.

- **CyclOT:** a bidirectional learned transport method with forward and backward maps, dynamic path matching, cycle consistency, and an annealed quadratic displacement penalty.
- **Entropic OT:** a Sinkhorn solver. For image experiments, we use the entropic coupling to obtain hard matches by selecting the target point with largest transport mass for each source point.
- **Neural OT:** a learned deterministic map trained through a neural optimal transport saddle objective.
- **ICNN OT:** an input-convex neural network baseline that learns convex Kantorovich potentials and defines the transport map as a potential gradient.
- **Rectified flow with reflow:** an iterative rectified-flow baseline that repeatedly retrains the velocity field using endpoint pairs induced by the previously learned flow, thereby straightening the learned transport trajectories.
- **Mini-batch OT-guided flow matching:** an OT-CFM-style baseline that first computes mini-batch Sinkhorn pairings and then trains a velocity field on the induced straight-line paths.

Evaluation Metrics To evaluate the numerical performance of comparison methods, we adopt the following metrics among the baseline datasets, when appropriate:

- **Empirical transport cost:** For image experiments, the reported cost is the average squared displacement between a source sample and its transported output. A lower value means that the learned map moves samples by a shorter distance. This quantity estimates the forward directional component of the quadratic action $\mathcal{A}$. *It is not, by itself, a measure of OT-map quality. In particular, a map that remains close to the source distribution can have low cost while failing to reach the target distribution. We therefore interpret transport cost together with endpoint-matching metrics.*
- **Target-domain rate and target probability score:** The target-domain rate (TDR) is the fraction of transported samples classified as belonging to the target domain, while the target probability score (TPS) is the average classifier probability assigned to the target domain. Higher values indicate stronger classifier-level target-domain transfer, which serves as a practical proxy of less mode collapse.
- **Feature-distribution discrepancy:** FID [38] compares transported and target samples through their empirical means and covariances in a pretrained feature space. A lower FID indicates closer agreement of these feature distributions. MMD [39] compares transported and target samples through kernel mean embeddings, with a lower empirical MMD² indicating closer distributional agreement in the chosen feature space and kernel family. Both quantities are finite-sample, representation-dependent proxies for marginal distribution matching performance.

- **Class coverage:** For MNIST, Class-TV measures the total variation distance between the predicted target-class proportions of transported samples and those of real target samples. A lower value indicates better coverage of digits 5, $\dots$, 9. Because Class-TV is computed conditionally on samples classified into the target domain, it must be interpreted jointly with TDR: a method may have balanced proportions among a small number of successful outputs while failing to transport most source samples.
- **Local mixing:** Enrich- k 100 compares local neighborhoods in the pooled transported and target sample sets. When the two sets have an equal size that is sufficiently large, a value near 0.5 indicates local mixing. We therefore report $|\text{Enrich}_{100} - 0.5|$, for which a smaller value is better. This measures local mixing in the chosen representation and complements global statistics such as FID and MMD.
- **Perturbation-signature correlation:** For the single-cell task, signature correlation measures agreement between the population-level feature shift produced by the learned map and the observed control-to-treatment feature shift. A higher value indicates better recovery of the average perturbation pattern.
- **Cycle consistency and map collapse:** Forward and reverse cycle errors measure how close is $g_\phi(f_\theta(x))$ and $f_\theta(g_\phi(y))$ to x and y on average, respectively. In the experiments below, the displayed cycle reconstructions provide a qualitative version of this diagnostic. They do not establish global injectivity or exclude collapse on unobserved or low-probability regions.

Computational measures. Training time measures wall-clock optimization cost under the stated hardware and software configuration. Image exposures count the number of source or target training examples accessed during optimization and therefore measure sample usage rather than transport quality. These quantities should be interpreted separately from statistical performance.

Connection with the theory. The empirical quantities mirror the theoretical roles of the three objective terms numerically with a goal to assess whether the observed tradeoffs are consistent with the mechanism underlying Theorem 4.4 and the structural results of Section 5. Distributional metrics probe endpoint agreement, empirical squared displacement probes the action term, and held-out cycle error probes approximate inverse compatibility. Their joint use is essential: low displacement without endpoint agreement can favor a near-identity map, while low marginal agreement or cycle error is insufficient in selecting the Brenier map without effectively penalizing the displacement loss.

6.2 Swiss roll: geometric fidelity in low dimensions

Setup. We first consider transport from a two-dimensional Gaussian source to a Swiss-roll target. This experiment isolates the geometric aspect of the problem: both endpoint distributions and the induced transport paths can be visualized directly in the data space. The goal is not only to match the target marginal, but also to learn a transport map whose pointwise displacements are coherent with the geometry of the target manifold.

Results. Figure 2 provides a qualitative visualization of the learned bidirectional transport. Marginal matching shows that the forward map transports Gaussian samples to the Swiss-roll distribution, while the backward map returns target samples to the Gaussian source. Cycle reconstructions approximately preserve target samples, indicating approximate inverse consistency

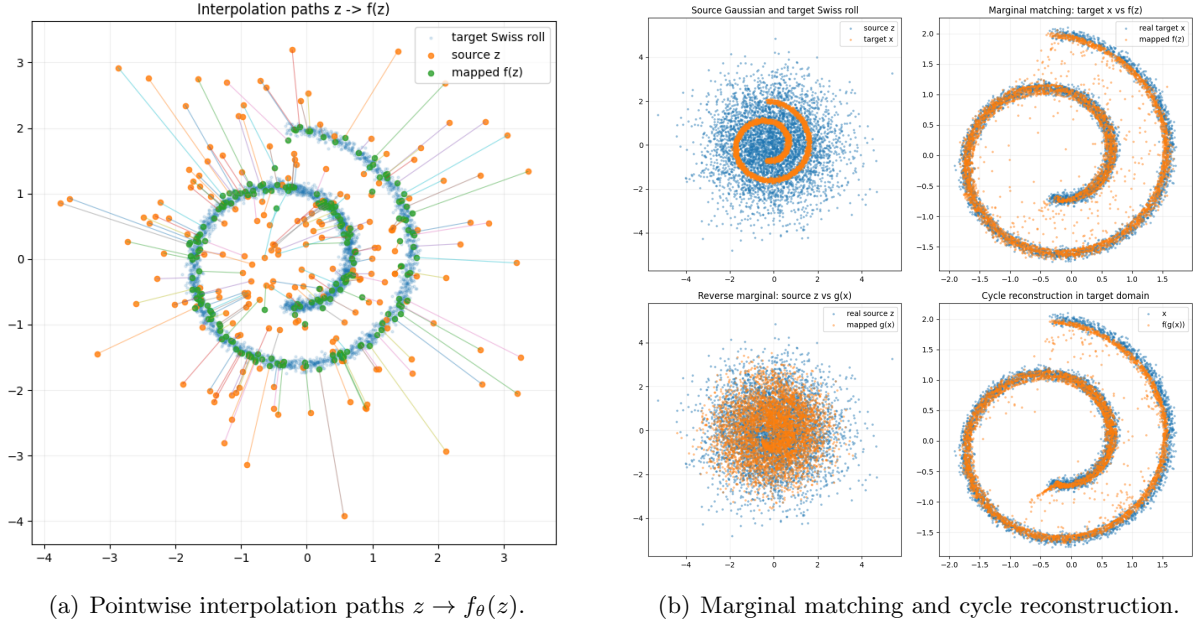


Figure 2: Qualitative visualization of the Gaussian-to-Swiss-roll experiment. The left panel shows pointwise map-induced interpolation paths from source samples to transported endpoints. The right panel shows forward marginal matching, reverse marginal matching, and target-domain cycle reconstruction.

between the two maps. The interpolation paths from source samples z to their endpoints $f_\theta(z)$ illustrate the displacement field induced by the learned forward map.

6.3 MNIST: transport from digits 0–4 to digits 5–9

We next study unpaired transport between two handwritten digit domains. The source distribution consists of MNIST [40] digits 0, 1, 2, 3, 4, and the target distribution consists of digits 5, 6, 7, 8, 9. This experiment tests whether a method can learn a global map that transports samples into the target digit domain while maintaining controlled displacement and avoiding trivial low-cost solutions that remain close to the source images.

Setup. All neural methods use MNIST images normalized to $[-1, 1]$. The proposed method, CyclOT, learns a forward map f_θ from digits 0–4 to digits 5–9 and a backward map g_ϕ from digits 5–9 to digits 0–4, both implemented as residual CNN encoder–decoder maps.

For entropic OT, we solve a finite-sample Sinkhorn problem between sampled source images and sampled target images using squared Euclidean pixel cost. Given the entropic coupling P , we form a hard matching by assigning each source image x_i to the target image $y_{j^*(i)}$, where

$$j^*(i) = \arg \max_j P_{ij}.$$

This baseline is transductive: it provides a strong finite-sample OT reference but does not learn a global map for unseen source images. For neural OT, ICNN OT, rectified flow, and mini-batch OT-guided flow matching, we train global maps or ODE flows and evaluate them on held-out source test images.

Method	Global map?	Cost ↓	TDR ↑	TPS ↑	FID ↓	Class-TV ↓	Train Time (min) ↓
Sinkhorn	No	245.201	0.992	0.984	93.773	0.072	–
CyclOT	Yes	212.43	0.957	0.954	92.269	0.039	7
Neural OT	Yes	133.68	0.956	0.952	120.45	0.130	198
ICNN OT	Yes	210.82	0.805	0.749	1565.65	0.162	43
Rectified flow	Yes	246.18	0.565	0.567	1282.97	0.171	29
OT-CFM	Yes	100.87	0.390	0.393	841.36	0.241	32

Table 1: MNIST 0–4 $\rightarrow$ 5–9 transport. We report transport cost, target-domain rate (TDR), target probability score (TPS), MNIST-feature FID, and total variation distance between predicted and real target-class distributions over digits 5, $\dots$, 9. Sinkhorn hard matching is a finite-sample transductive baseline and does not define a learned global map. Among learned global maps, CyclOT achieves the best MNIST-feature FID and target class-distribution matching, while maintaining high target-domain validity.



Figure 3: MNIST 0–4 $\rightarrow$ 5–9 transport.

Evaluation. Across experiments, we report metrics that separately quantify transport structure, target-population matching, and computational efficiency. For the MNIST transport task, these include the per-sample squared transport cost, target-domain rate (TDR), target probability score (TPS), MNIST-feature FID, class-distribution total variation distance (Class-TV), and training time. Detailed definitions and computation procedures are provided in Appendix C.1.

Results. Table 1 reports the quantitative results. Among learned global maps, CyclOT achieves the best overall tradeoff between target-domain generation quality and transport structure. It attains a target-domain rate of 0.957, indicating that most transported images are classified as digits 5–9, while also achieving the lowest MNIST-feature FID among learned methods. The class-distribution TV distance is also the smallest among learned methods, suggesting that the transported samples better cover the target digit classes rather than collapsing to a small subset of target digits.

Neural OT achieves a lower transport cost but worse FID and class-distribution discrepancy, suggesting that lower displacement alone does not necessarily imply better target-distribution matching. The Sinkhorn baseline achieves high target-domain validity by directly matching to observed target samples, but does not define a learned global transport map.

Figure 3 presents qualitative results on MNIST. The first row shows the source digits (0–4), while the second row shows their corresponding mapped digits in the target classes (5–9). Additional visualization results are provided in Appendix D.1. The visual results are consistent with the quantitative metrics: CyclOT produces samples that are visually closer to target digits 5–9, while the lower-cost flow baselines tend to preserve too much of the original 0–4 digit structure.

6.4 CelebA: bidirectional transport on natural images

We then consider a transport task on CelebA [41], mapping between female and male face distributions. Compared with MNIST, this setting is substantially more challenging: the data lie on a richer natural-image manifold, and successful transport must preserve shared facial structure while modifying domain-level attributes.

Setup. We partition CelebA using the **Male** attribute, taking female faces as the source domain and male faces as the target domain. The forward map f_θ learns to transport female faces to the male domain, while the backward map g_ϕ learns the reverse mapping from male faces to the female domain. Images are center-cropped, resized to 64×64 , and normalized to $[-1, 1]$. The forward and backward maps are implemented as residual encoder-decoder CNNs with skip connections.

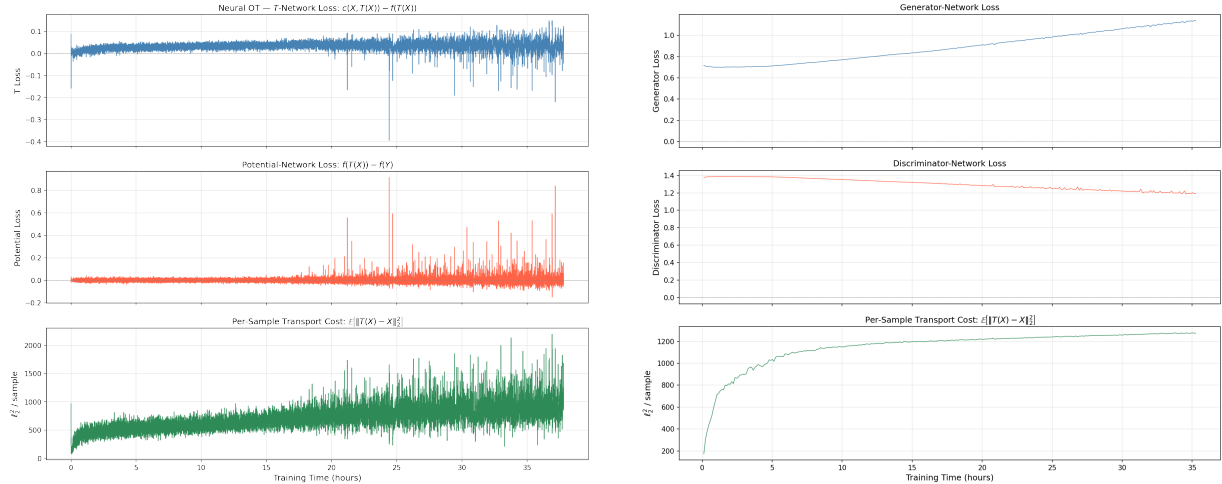
Evaluation. Across all experiments, we present qualitative results and quantitatively compare CyclOT with Neural OT and ICNN OT using seven evaluation metrics: transport cost, Fréchet Inception Distance (FID), target domain rate(TDR), Target Probability Score(TPS), Enrich-k100 and Female Image Exposures(FIE), and Male Image Exposures(MIE). The results are reported in Table 2. The TDR measures the proportion of transported source-female images that are classified as male, while the TPS measures the average predicted probability assigned to the male class. Both metrics are computed using the FairFace classifier [42]. As a real-data reference, we also report FairFace results on real CelebA female and male images in Table 3. Female/Male Image Exposures denote the total number of times female/male images from the CelebA training set are accessed during training. Detailed definitions and computation procedures of each metric are provided in Appendix C.1.

Results. Figure 5 qualitatively demonstrates the learned bidirectional transport on CelebA. The forward and backward maps modify gender-related visual attributes while largely preserving pose, facial layout, and background. The generated images do not exhibit obvious visual artifacts, suggesting that the learned maps remain stable under cross-domain translation. The cycle reconstructions $g_\phi(f_\theta(x))$ and $f_\theta(g_\phi(y))$ are visually close to the original samples, indicating that the cycle-consistency term helps retain sample-level information and discourages many-to-one collapse. Although some fine details are blurred or slightly smoothed, the reconstructions preserve the main semantic and structural content of the inputs. Additional experimental results are presented in Appendix D.2.

Table 2 shows that Neural OT achieves a lower transport cost, whereas CyclOT achieves higher TDR and TPS. Moreover, our results are closer to the FairFace baseline reported in Table 3: The TDR is 92.0% versus baseline 92.6%, and the TPS is 0.8602 ± 0.20 versus baseline 0.8893 ± 0.21 . In addition, Figure 7 further shows that CyclOT reaches a higher TDR more rapidly during training.

Figure 6 shows that, under the same training budget, CyclOT produces stronger and more consistent female-to-male translations than Neural OT while preserving pose, background, and coarse facial structure. In contrast, Neural OT often retains more source-domain appearance, despite achieving a lower transport cost. This comparison highlights that transport cost alone is not sufficient for learning an effective cross-domain map: when the target marginal is not enforced strongly enough, minimizing displacement can favor mappings that remain close to the source distribution. CyclOT instead achieves stronger target-domain alignment while preserving the overall structure of the input samples.

Figure 4(a) shows that Neural OT becomes unstable as training progresses and deteriorates after reaching an intermediate optimum. In contrast, Figure 4(b) shows more stable training dynamics.



(a) Training dynamics of Neural OT on the CelebA female-to-male translation task. Here, T denotes the learned transport map, and f denotes the learned Kantorovich potential.

(b) Training dynamics of CyclOT on the CelebA female-to-male translation task.

Figure 4: Training dynamics on the CelebA female-to-male translation task for Neural OT and CyclOT.

Method	Cost ↓	FID ↓	TDR ↑	TPS ↑	Enrich-k100 − 0.5 ↓	FIE ↓	MIE ↓
ICNN OT	2266.7034	305.2227	46.0%	0.4851 ± 0.25	0.3759	3.2M	35.2M
Neural OT	687.9927	17.8291	79.9%	0.7447 ± 0.27	0.0228	17.6M	1.6M
CyclOT	1069.7815	25.4478	92.0%	0.8602 ± 0.20	0.0357	2.73M	2.73M

Table 2: Quantitative comparison between Neural OT and CyclOT on the CelebA female-to-male translation task.

Under our BCE-type GAN objective, the balanced discriminator equilibrium corresponds to a generator loss of $\log 2$ and a discriminator loss of $2 \log 2$ under the summed-BCE convention. In addition, the transport cost curve shows that CyclOT converges smoothly during training.

6.5 Single-cell data: transport on scientific manifolds

Setup. We evaluate the methods on a single-cell 4i dataset from Bunne et al. [3] using the trametinib perturbation experiment. We take cells from the control condition as the source domain and trametinib-treated cells as the target domain. Each cell is represented by the 48 molecular features, and no cell-level correspondence between control and treated cells is assumed. The goal is therefore to learn an unpaired transport map from the control population to the trametinib-treated population. The forward and backward maps are implemented as residual MLPs in the single-cell feature space.

Evaluation. We evaluate the single-cell transport task using metrics that capture both transport regularity and target-population matching. Specifically, we report the *transport cost*, *MMD loss*, and *Enrich-k100*. Detailed definitions and implementation details for these metrics are provided in the appendix C.1.

Input set	TPS	Predicted male	Predicted female
Real source female	0.1278 ± 0.23	479 (9.6%)	4521 (90.4%)
Real target male	0.8893 ± 0.21	4629 (92.6%)	371 (7.4%)

Table 3: FairFace classifier baseline results on real CelebA source-female and target-male images. Each evaluation set contains 5,000 images. The classifier correctly predicts 90.4% of the source-female images as female and 92.6% of the real target-male images as male.

Method	Transport Cost ↓	MMD ↓	Enrich-k100 − 0.5 ↓	Signature Corr ↑
Sinkhorn	5.2320	–	–	–
CyclOT	2.8352	0.002472	0.00729	0.99093
Neural OT	3.1925	0.002879	0.01835	0.98936
ICNN OT	3.2515	0.002623	0.00456	0.98698
Rectified flow	9.4423	0.004876	0.02044	0.94781
OT-CFM	7.0150	0.012698	0.11412	0.86522

Table 4: 4i trametinib single-cell perturbation transport.

Results. Table 4 summarizes the control-to-trametinib transport results on the 4i dataset. Among the learned methods, CyclOT achieves the lowest transport cost, the lowest MMD, and the highest signature correlation, while also attaining competitive local neighborhood mixing. ICNN OT achieves the best Enrich-k100 score. Rectified flow and OT-guided flow matching exhibit substantially larger transport costs together with weaker target-population matching. Sinkhorn hard matching is included as a finite-sample transductive reference and does not define a learned global map.

6.6 Chest X-ray translation: transport on medical images

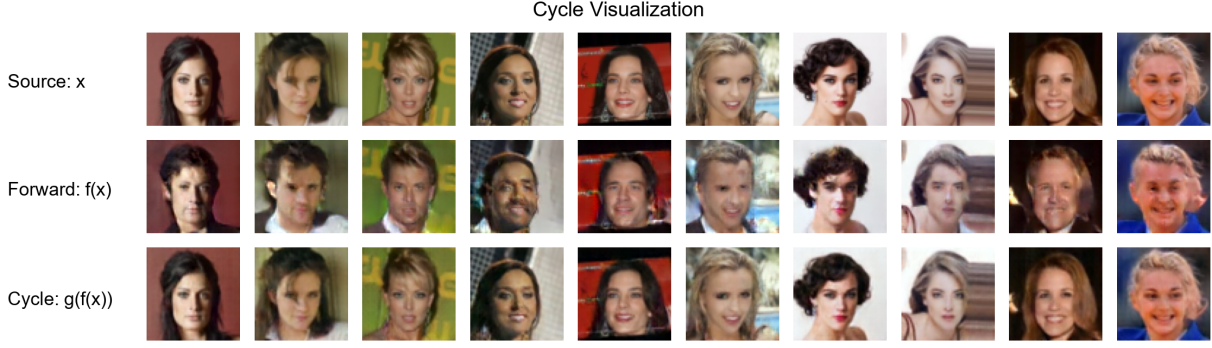
Method	Cost ↓	Enrich-k100 − 0.5 ↓	TPS ↑	Training Time (hours:minutes) ↓
Sinkhorn	5643.2637	–	–	<0:01
ICNN OT	65095.5671	0.4984	0.1513 ± 0.0409	1:01
Neural OT	1984.1612	0.3246	0.4762 ± 0.1860	19:15
CyclOT	1717.2522	0.3169	0.4098 ± 0.1873	8:41

Table 5: Quantitative comparison between ICNN OT, Neural OT, and CyclOT on the chest x-ray healthy-to-pneumonia transition task.

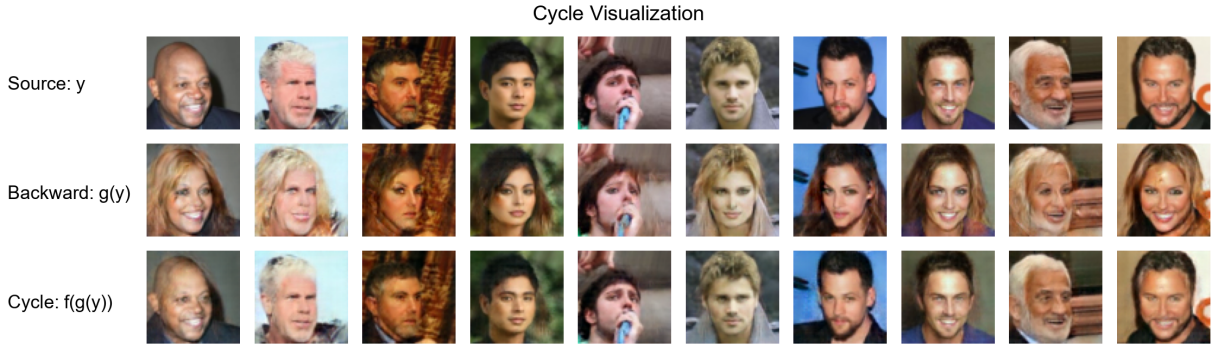
Setup. We further evaluate CyclOT on chest x-ray images to test whether the proposed forward-backward transport framework can operate on high-dimensional medical data. Compared with natural image attribute translation, chest radiographs impose a stronger structure-preservation requirement: a meaningful transport map should modify the domain-specific radiographic appearance while preserving the underlying anatomical layout, including the lung fields, ribs, clavicles, and global patient pose.

We use data from the RSNA Pneumonia Detection Challenge¹ for training [43]. The dataset

¹Data is available to download here: <https://nihcc.app.box.com/v/ChestXray-NIHCC>



(a) Female-to-male-to-female cycle.



(b) Male-to-female-to-male cycle.

Figure 5: Cycle visualizations on CelebA. The first block shows source female samples x , forward mapped samples $f_\theta(x)$, and cycle reconstructions $g_\phi(f_\theta(x))$. The second block shows source male samples y , backward mapped samples $g_\phi(y)$, and cycle reconstructions $f_\theta(g_\phi(y))$.

includes a total of 7,106 chest x-rays from patients with pneumonia and 9,790 x-rays from patients without pneumonia. We resize all images to 256×256 and use a random subset of 3,000 images from each class for training, and repeat the training three times for each method. The full dataset is used for evaluation.

Evaluation. To evaluate the utility of our and other methods at converting x-ray images between the two classes we use a classification model developed by Weng et al. [44] on the dataset provided by Rajpurkar et al. [45]. We report the transport cost, Enrich-k100, TPS, and total training time of each method. The results in Table 5 show the mapping from healthy images to ones with pneumonia. The results for the mapping in the other direction can be found in D.3.

Results. Figure 8 shows representative qualitative results. The first row shows source chest X-ray images x , the second row shows transported images $f_\theta(x)$, and the third row shows cycle reconstructions $g_\phi(f_\theta(x))$. The transported samples exhibit visible changes in global radiographic appearance, including contrast, opacity, and local texture patterns, while largely preserving the coarse anatomical geometry of the input images. In particular, the position of the thoracic cavity, lung fields, shoulders, and mediastinal structure remains aligned with the corresponding source image. The cycle reconstructions recover the main image structure and much of the source-domain

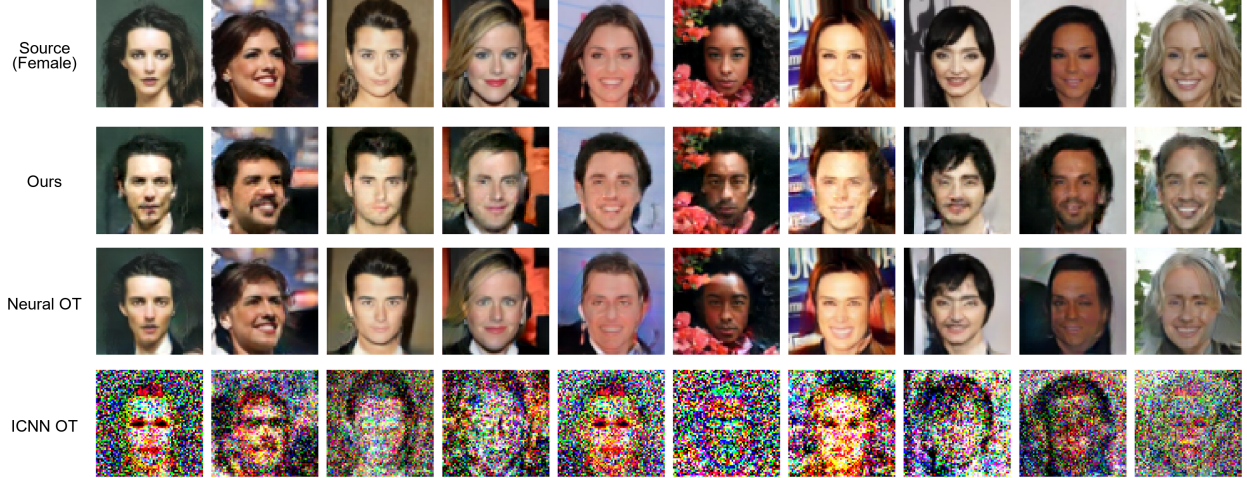


Figure 6: Comparison results for female-to-male face translation on CelebA. All three methods are trained for 4 hours under the same experimental setting.

Input set	TPS
Real source healthy	0.1876 ± 0.1382
Real target pneumonia	0.6254 ± 0.1902

Table 6: Pneumonia classifier baseline results on real source (healthy) and target (pneumonia) images.

appearance, suggesting that the learned map is not simply collapsing different inputs to a small set of target-like images.

6.7 Takeaway

Taken together, the experiments illustrate a consistent pattern. Entropic OT remains a strong classical baseline for finite-sample marginal alignment, especially in low dimension, but it does not directly provide a learned global map and may require additional projection or regression steps. Neural OT provides expressive learned maps, but without explicit reverse consistency it may admit transport solutions that match target-domain labels while failing to match the full target distribution. ICNN OT encodes the Brenier structure but is difficult to optimize directly in high-dimensional pixel space. Flow-matching baselines provide flexible ODE transports, but without sufficiently informative couplings they may favor low-displacement maps that do not fully reach the target domain. CyclOT combines the flexibility of learned transport with bidirectional structural regularization, leading to maps that are not only distributionally accurate but also more stable and coherent along intermediate paths.

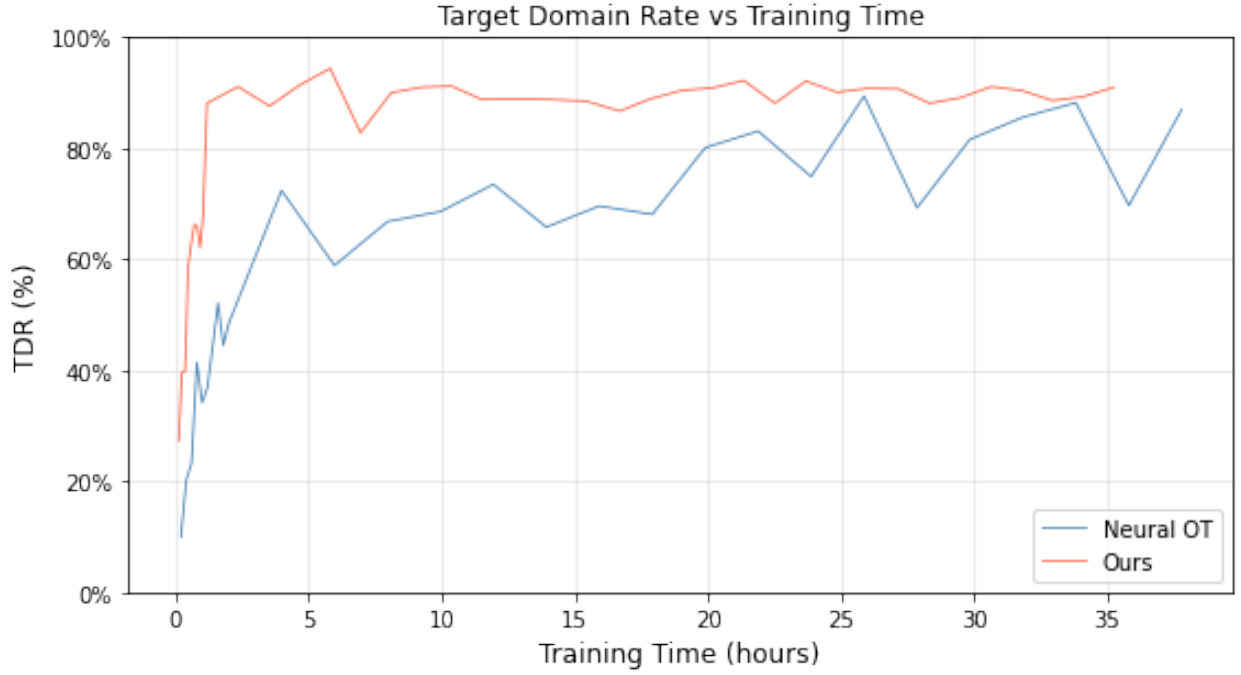


Figure 7: Target Domain Rate (TDR) versus training time for Neural OT and CyclOT.

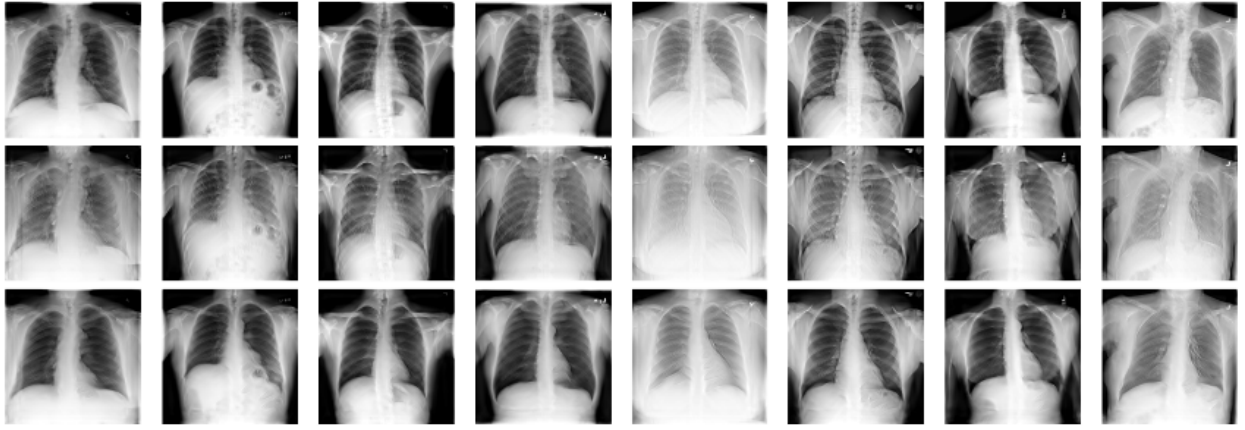
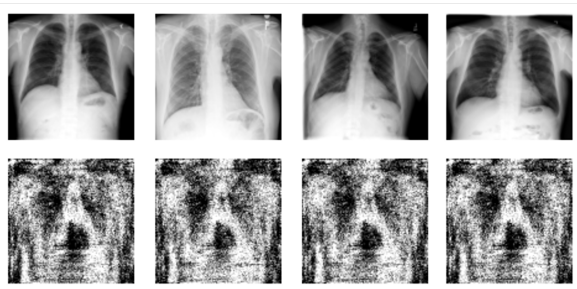
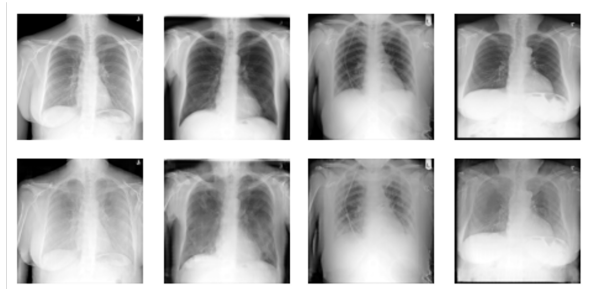


Figure 8: Chest x-ray transport results for CyclOT. The first row shows source (healthy) radio-graphs x , the second row shows the transported samples $f_{\theta}(x)$, and the third row shows the cycle reconstructions $g_{\phi}(f_{\theta}(x))$. The learned transport changes the target-domain radiographic appearance while preserving the coarse anatomical layout.



(a) Images created by a model trained with ICNN OT.



(b) Images created by a model trained with Neural OT.

Figure 9: Chest x-ray transport results for other tested methods. The first row shows source (healthy) radiographs x and the second row shows the transported samples $f_{\theta}(x)$.

References

- [1] Jun-Yan Zhu, Taesung Park, Phillip Isola, and Alexei A. Efros. Unpaired image-to-image translation using cycle-consistent adversarial networks. In *Proceedings of the IEEE International Conference on Computer Vision*, pages 2223–2232, 2017.
- [2] Nicolas Courty, Rémi Flamary, Devis Tuia, and Alain Rakotomamonjy. Optimal transport for domain adaptation. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 39(9):1853–1865, 2017.
- [3] Charlotte Bunne, Stefan G. Stark, Gabriele Gut, Jacobo Sarabia Del Castillo, Mitchell Levesque, Kjong-Van Lehmann, Lucas Pelkmans, Andreas Krause, and Gunnar Rätsch. Learning single-cell perturbation responses using neural optimal transport. *Nature Methods*, 20:1759–1768, 2023.
- [4] Xingchao Liu, Chengyue Gong, and Qiang Liu. Flow straight and fast: Learning to generate and transfer data with rectified flow. In *International Conference on Learning Representations*, 2023. arXiv:2209.03003.
- [5] Cédric Villani. *Optimal Transport: Old and New*. Springer, 2009.
- [6] Filippo Santambrogio. *Optimal Transport for Applied Mathematicians*. Birkhäuser, 2015.
- [7] Gabriel Peyré and Marco Cuturi. Computational optimal transport. *Foundations and Trends in Machine Learning*, 11(5–6):355–607, 2019.
- [8] Yann Brenier. Polar factorization and monotone rearrangement of vector-valued functions. *Communications on Pure and Applied Mathematics*, 44(4):375–417, 1991.
- [9] Robert J. McCann. A convexity principle for interacting gases. *Advances in Mathematics*, 128(1):153–179, 1997.
- [10] Jean-David Benamou and Yann Brenier. A computational fluid mechanics solution to the monge–kantorovich mass transfer problem. *Numerische Mathematik*, 84:375–393, 2000.
- [11] Marco Cuturi. Sinkhorn distances: Lightspeed computation of optimal transport. In *Advances in Neural Information Processing Systems*, volume 26, 2013.
- [12] Aram-Alexandre Pooladian, Heli Ben-Hamu, Carles Domingo-Enrich, Brandon Amos, Yaron Lipman, and Ricky T. Q. Chen. Multisample flow matching: Straightening flows with minibatch couplings. *arXiv preprint arXiv:2304.14772*, 2023.
- [13] Alexander Tong, Kilian Fatras, Nikolay Malkin, Guillaume Huguet, Yanlei Zhang, Jarrid Rector-Brooks, Guy Wolf, and Yoshua Bengio. Improving and generalizing flow-based generative models with minibatch optimal transport. *arXiv preprint arXiv:2302.00482*, 2023.
- [14] Kilian Fatras, Younes Zine, Rémi Flamary, Rémi Gribonval, and Nicolas Courty. Learning with minibatch Wasserstein: Asymptotic and gradient properties. *arXiv preprint arXiv:1910.04091*, 2019.
- [15] Martin Arjovsky, Soumith Chintala, and Léon Bottou. Wasserstein generative adversarial networks. In *Proceedings of the 34th International Conference on Machine Learning*, volume 70 of *Proceedings of Machine Learning Research*, pages 214–223. PMLR, 2017.

- [16] Alexander Korotin, Daniil Selikhanovych, and Evgeny Burnaev. Neural optimal transport. In *International Conference on Learning Representations*, 2023.
- [17] Amirhossein Taghvaei and Amin Jalali. 2-Wasserstein approximation via restricted convex potentials with application to improved training for gans. *arXiv preprint arXiv:1902.07197*, 2019.
- [18] Ashok Vardhan Makkuva, Amirhossein Taghvaei, Sewoong Oh, and Jason D. Lee. Optimal transport mapping via input convex neural networks. In *Proceedings of the 37th International Conference on Machine Learning*, pages 6672–6681, 2020.
- [19] Alexander Korotin, Vage Egiazarian, Arip Asadulaev, Alexander Safin, and Evgeny Burnaev. Wasserstein-2 generative networks. *arXiv preprint arXiv:1909.13082*, 2019.
- [20] Brandon Amos, Lei Xu, and J. Zico Kolter. Input convex neural networks. In *Proceedings of the 34th International Conference on Machine Learning*, volume 70 of *Proceedings of Machine Learning Research*, pages 146–155. PMLR, 2017.
- [21] Nikita Kornilov, Petr Mokrov, Alexander Gasnikov, and Alexander Korotin. Optimal flow matching: Learning straight trajectories in just one step. In *Advances in Neural Information Processing Systems*, 2024. *arXiv:2403.13117*.
- [22] Alireza Mousavi-Hosseini, Stephen Y. Zhang, Michal Klein, and Marco Cuturi. Flow matching with semidiscrete couplings. *arXiv preprint arXiv:2509.25519*, 2025.
- [23] Ling kai Kong, Molei Tao, Yang Liu, Bryan Wang, Jinmiao Fu, Chien-Chih Wang, and Huidong Liu. Alignflow: Improving flow-based generative models with semi-discrete optimal transport. In *International Conference on Learning Representations*, 2026.
- [24] Yichi Zhang, Yici Yan, Alex Schwing, and Zhizhen Zhao. Hierarchical rectified flow matching with mini-batch couplings. *arXiv preprint arXiv:2507.13350*, 2025.
- [25] Chenrui Ma, Xi Xiao, Tianyang Wang, Xiao Wang, and Yanning Shen. Stochastic interpolants via conditional dependent coupling, 2026. OpenReview preprint.
- [26] Shizhou Xu and Thomas Strohmer. Fair data representation for machine learning at the pareto frontier. *Journal of Machine Learning Research*, 24(331):1–63, 2023.
- [27] Yaron Lipman, Ricky T. Q. Chen, Heli Ben-Hamu, Maximilian Nickel, and Matt Le. Flow matching for generative modeling. *arXiv preprint arXiv:2210.02747*, 2022.
- [28] Yang Song, Prafulla Dhariwal, Mark Chen, and Ilya Sutskever. Consistency models. *arXiv preprint arXiv:2303.01469*, 2023.
- [29] Dongjun Kim, Chieh-Hsin Lai, Wei-Hsiang Liao, Naoki Murata, Yuhta Takida, Toshimitsu Ue-saka, Yutong He, Yuki Mitsufuji, and Stefano Ermon. Consistency trajectory models: Learning probability flow ode trajectory of diffusion. *arXiv preprint arXiv:2310.02279*, 2023.
- [30] Tim Salimans and Jonathan Ho. Progressive distillation for fast sampling of diffusion models. *arXiv preprint arXiv:2202.00512*, 2022.
- [31] Kevin Frans, Danijar Hafner, Sergey Levine, and Pieter Abbeel. One step diffusion via shortcut models. *arXiv preprint arXiv:2410.12557*, 2024.

- [32] Zhengyang Geng, Mingyang Deng, Xingjian Bai, J. Zico Kolter, and Kaiming He. Mean flows for one-step generative modeling. *arXiv preprint arXiv:2505.13447*, 2025.
- [33] Nicholas Matthew Boffi, Michael Samuel Albergo, and Eric Vanden-Eijnden. Flow map matching with stochastic interpolants: A mathematical framework for consistency models. *Transactions on Machine Learning Research*, 2025.
- [34] Afonso S Bandeira, Amit Singer, and Thomas Strohmer. Mathematics of Data Science. *arXiv e-prints*, *arxiv:2607.11938*, 2026.
- [35] Kurt Hornik. Approximation capabilities of multilayer feedforward networks. *Neural Networks*, 4(2):251–257, 1991.
- [36] Moshe Leshno, Vladimir Ya. Lin, Allan Pinkus, and Shimon Schocken. Multilayer feedforward networks with a nonpolynomial activation function can approximate any function. *Neural Networks*, 6(6):861–867, 1993.
- [37] George Cybenko. Approximation by superpositions of a sigmoidal function. *Mathematics of Control, Signals, and Systems*, 2(4):303–314, 1989.
- [38] Martin Heusel, Hubert Ramsauer, Thomas Unterthiner, Bernhard Nessler, and Sepp Hochreiter. Gans trained by a two time-scale update rule converge to a local nash equilibrium. In *Advances in Neural Information Processing Systems*, volume 30, 2017.
- [39] Arthur Gretton, Karsten M. Borgwardt, Malte J. Rasch, Bernhard Schölkopf, and Alexander Smola. A kernel two-sample test. *Journal of Machine Learning Research*, 13(25):723–773, 2012.
- [40] Y. Lecun, L. Bottou, Y. Bengio, and P. Haffner. Gradient-based learning applied to document recognition. *Proceedings of the IEEE*, 86(11):2278–2324, 1998.
- [41] Ziwei Liu, Ping Luo, Xiaogang Wang, and Xiaoou Tang. Deep learning face attributes in the wild. In *Proceedings of the IEEE International Conference on Computer Vision*, pages 3730–3738, 2015.
- [42] Kimmo Kärkkäinen and Jungseock Joo. Fairface: Face attribute dataset for balanced race, gender, and age. *CoRR*, abs/1908.04913, 2019.
- [43] Xiaosong Wang, Yifan Peng, Le Lu, Zhiyong Lu, Mohammadhadi Bagheri, and Ronald M Summers. Chestx-ray8: Hospital-scale chest x-ray database and benchmarks on weakly-supervised classification and localization of common thorax diseases. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pages 2097–2106, 2017.
- [44] Xinyu Weng, Nan Zhuang, Jingjing Tian, and Yingcheng Liu. Chexnet for classification and localization of thoracic diseases. GitHub, 2017.
- [45] Pranav Rajpurkar, Jeremy Irvin, Kaylie Zhu, Brandon Yang, Hershel Mehta, Tony Duan, Daisy Yi Ding, Aarti Bagul, Curtis P. Langlotz, Katie S. Shpanskaya, Matthew P. Lungren, and Andrew Y. Ng. Chexnet: Radiologist-level pneumonia detection on chest x-rays with deep learning. *CoRR*, abs/1711.05225, 2017.
- [46] Maximilian Seitzer. pytorch-fid: FID Score for PyTorch. <https://github.com/mseitzer/pytorch-fid>, August 2020. Version 0.3.0.

A Background & Preliminaries

A.1 Flow matching & rectified flow

Flow matching learns a time-dependent vector field whose induced probability flow transports a source distribution to a target distribution [27]. Specifically, one considers the ODE

$$\frac{dZ_t}{dt} = v_\theta(Z_t, t), \quad Z_0 \sim \mu_0,$$

and trains v_θ using a prescribed reference interpolation between samples from μ_0 and μ_1 . A common choice is the linear interpolation

$$X_t = (1 - t)X_0 + tX_1, \quad X_0 \sim \mu_0, \quad X_1 \sim \mu_1,$$

where (X_0, X_1) is sampled from a chosen coupling, often the independent coupling. The corresponding reference velocity is

$$u(X_0, X_1) = X_1 - X_0.$$

Flow matching then minimizes the regression objective

$$\mathcal{L}_{\text{FM}}(\theta) = \int_0^1 \mathbb{E} \left[\|X_1 - X_0 - v_\theta(X_t, t)\|^2 \right] dt.$$

At the population level, the minimizer over all measurable vector fields is

$$v^*(z, t) = \mathbb{E}[X_1 - X_0 \mid X_t = z].$$

The reference marginals $\mu_t = \mathcal{L}(X_t)$ then satisfy the continuity equation

$$\partial_t \mu_t + \nabla \cdot (\mu_t v_t^*) = 0$$

in the weak sense. If this continuity equation is well posed for v^* , then the ODE driven by v^* has the same time marginals as the reference interpolation.

Rectified flow applies this projection iteratively. After learning a velocity field, it constructs new endpoint pairings from the learned flow and retrain on the resulting trajectories. This reflow procedure preserves the endpoint marginals while progressively straightening the learned paths.

A.2 Independence constraints & Wasserstein barycenters.

Following Xu and Strohmer [26], we recall the connection between independence-constrained L^2 projection and Wasserstein barycenters. Let $Z \in \{0, 1\}$ be a group label with

$$\mathbb{P}(Z = k) = \lambda_k, \quad \mu_k = \mathcal{L}(X \mid Z = k), \quad k \in \{0, 1\}.$$

Consider the independence-constrained projection problem

$$\inf_{\bar{X}: \bar{X} \perp Z} \mathbb{E} [\|X - \bar{X}\|^2].$$

This constraint means that the conditional law of $\bar{X}$ is the same for both groups. Equivalently, there exists a common distribution ν such that

$$\mathcal{L}(\bar{X} \mid Z = 0) = \mathcal{L}(\bar{X} \mid Z = 1) = \nu.$$

For a fixed common law ν , the optimal way to map the k -th group distribution μ_k to ν under quadratic distortion has cost

$$W_2^2(\mu_k, \nu).$$

Therefore, the independence-constrained projection problem reduces to choosing the best common law:

$$\inf_{\bar{X}: \bar{X} \perp Z} \mathbb{E} \|X - \bar{X}\|^2 = \inf_{\nu \in \mathcal{P}_2(\mathbb{R}^d)} \sum_{k=0}^1 \lambda_k W_2^2(\mu_k, \nu).$$

Any minimizer ν^* of the right-hand side is the quadratic Wasserstein barycenter of the group-conditional distributions [26].

This result shows that an independence constraint can convert an L^2 projection problem into a Wasserstein barycenter problem: independence enforces a shared output law, while the quadratic projection cost measures the optimal transport cost from each group distribution to that shared law.

B Proofs

B.1 Proof of Lemma 4.1

Proof. Since $\mathcal{X}$ is nonempty, closed, and convex, its Euclidean metric projection

$$\Pi_{\mathcal{X}} : \mathbb{R}^d \rightarrow \mathcal{X}$$

is well defined and 1-Lipschitz. Let

$$E_0 := \{x \in \mathcal{X} : S(T(x)) = x\}.$$

Then $\mu_0(E_0) = 1$. Fix $0 < \eta < 1/3$. By Lusin's theorem, there exist compact sets $C_0, C_1 \subset \mathcal{X}$ such that $\mu_0(C_0) > 1 - \eta$ and $\mu_1(C_1) > 1 - \eta$ with $T|_{C_0}$ and $S|_{C_1}$ are continuous. Since $T_{\#}\mu_0 = \mu_1$,

$$\mu_0(T^{-1}(C_1)) = \mu_1(C_1) > 1 - \eta.$$

Consequently,

$$\mu_0(C_0 \cap E_0 \cap T^{-1}(C_1)) > 1 - 2\eta.$$

By inner regularity, we can choose a compact set $K_0 \subset C_0 \cap E_0 \cap T^{-1}(C_1)$ such that

$$\mu_0(K_0) > 1 - 3\eta,$$

and define $K_1 := T(K_0)$. Because $T|_{K_0}$ is continuous, K_1 is compact. Also, since $K_0 \subset T^{-1}(C_1)$, we have $K_1 \subset C_1$. Moreover,

$$\mu_1(K_1) = \mu_0(T^{-1}(K_1)) \geq \mu_0(K_0) > 1 - 3\eta.$$

Since $K_0 \subset E_0$, $S \circ T = \text{Id}$ on K_0 . If $y \in K_1$, then $y = T(x)$ for some $x \in K_0$, and hence

$$T(S(y)) = T(S(T(x))) = T(x) = y.$$

Thus, $T \circ S = \text{Id}$ on K_1 . Applying the Tietze extension theorem coordinatewise to obtain continuous maps

$$A_\eta, B_\eta : \mathcal{X} \rightarrow \mathbb{R}^d$$

such that $A_\eta = T$ on K_0 and $B_\eta = S$ on K_1 . Define

$$\tilde{T}_\eta := \Pi_{\mathcal{X}} \circ A_\eta, \quad \tilde{S}_\eta := \Pi_{\mathcal{X}} \circ B_\eta.$$

Then $\tilde{T}_\eta, \tilde{S}_\eta \in C(\mathcal{X}; \mathcal{X})$ and the projection fixes the prescribed values, so $\tilde{T}_\eta = T$ on K_0 and $\tilde{S}_\eta = S$ on K_1 . By uniform density, for every $e > 0$, there exist $T_{\eta,e} \in \mathcal{F}$ and $S_{\eta,e} \in \mathcal{G}$ such that

$$\|T_{\eta,e} - \tilde{T}_\eta\|_\infty < e, \quad \|S_{\eta,e} - \tilde{S}_\eta\|_\infty < e.$$

All these maps are $\mathcal{X}$ -valued, so their compositions remain in $\mathcal{X}$. For $H \in C(\mathcal{X}; \mathbb{R}^d)$, let

$$\omega_H(t) := \sup\{\|H(u) - H(v)\| : u, v \in \mathcal{X}, \|u - v\| \leq t\}.$$

Compactness of $\mathcal{X}$ implies $\omega_H(t) \rightarrow 0$ as $t \downarrow 0$. For $x \in K_0$,

$$\begin{aligned} \|S_{\eta,e}(T_{\eta,e}(x)) - x\| &\leq \|S_{\eta,e}(T_{\eta,e}(x)) - \tilde{S}_\eta(T_{\eta,e}(x))\| \\ &\quad + \|\tilde{S}_\eta(T_{\eta,e}(x)) - \tilde{S}_\eta(T(x))\| \\ &\quad + \|\tilde{S}_\eta(T(x)) - x\| \\ &\leq e + \omega_{\tilde{S}_\eta}(e) + 0. \end{aligned}$$

The final term vanishes because $T(x) \in K_1$. Similarly, for $y \in K_1$, we obtain

$$\|T_{\eta,e}(S_{\eta,e}(y)) - y\| \leq e + \omega_{\tilde{T}_\eta}(e).$$

Set $D_{\mathcal{X}} := \text{diam}(\mathcal{X})$. On the complementary sets with measure at most 3η , all points and compositions lie in $\mathcal{X}$, so the corresponding errors are bounded by $D_{\mathcal{X}}$. It follows that

$$\mathcal{C}(T_{\eta,e}, S_{\eta,e}) \leq (e + \omega_{\tilde{S}_\eta}(e))^2 + (e + \omega_{\tilde{T}_\eta}(e))^2 + 6D_{\mathcal{X}}^2\eta,$$

and

$$\|T_{\eta,e} - T\|_{L^2(\mu_0)}^2 \leq e^2 + 3D_{\mathcal{X}}^2\eta, \quad \|S_{\eta,e} - S\|_{L^2(\mu_1)}^2 \leq e^2 + 3D_{\mathcal{X}}^2\eta.$$

Choose $\eta_j \downarrow 0$ with $0 < \eta_j < 1/3$. For each j , perform the preceding construction with $\eta = \eta_j$, and then choose $e_j > 0$ so that

$$e_j < \frac{1}{j}, \quad e_j + \omega_{\tilde{T}_{\eta_j}}(e_j) < \frac{1}{j}, \quad e_j + \omega_{\tilde{S}_{\eta_j}}(e_j) < \frac{1}{j}.$$

Setting $(T_j, S_j) := (T_{\eta_j, e_j}, S_{\eta_j, e_j})$ and using the preceding estimates proves both conclusions. We are done. $\square$

B.2 Proof of Appendix 5.1

Proof. Markov's inequality gives

$$\mathbb{P}_{X \sim \mu_0}(\|g(f(X)) - X\| > \varepsilon) \leq \frac{1}{\varepsilon^2} \int_{\mathcal{X}} \|g(f(x)) - x\|^2 d\mu_0(x)$$

and

$$\mathbb{P}_{Y \sim \mu_1}(\|f(g(Y)) - Y\| > \varepsilon) \leq \frac{1}{\varepsilon^2} \int_{\mathcal{X}} \|f(g(y)) - y\|^2 d\mu_1(y).$$

Adding these inequalities yields

$$p_\varepsilon(f, g) \leq \frac{\mathcal{C}(f, g)}{\varepsilon^2}.$$

The bound by 2 follows because $p_\varepsilon(f, g)$ is the sum of two probabilities. $\square$

C Evaluation Metrics & Experimental Details

C.1 Evaluation metrics

C.1.1 Transport cost evaluation

Transport Cost Transport cost in task MNIST, CelebA and Chest X-ray uses average per-sample squared L_2 distance in pixel space, computed as the sum of the squared pixel-wise differences between each source image and its transported image, averaged over all samples.

$$\text{Transport Cost} = \frac{1}{N} \sum_{i=1}^N \|T(x_i) - x_i\|_2^2 = \frac{1}{N} \sum_{i=1}^N \sum_{c=1}^C \sum_{h=1}^H \sum_{w=1}^W (T(x_i)_{c,h,w} - x_{i,c,h,w})^2.$$

- N denotes the total number of images used for evaluation;
- C denotes the number of image channels;
- H denotes the image height;
- W denotes the image width.

C.1.2 Marginal matching quality

Target-Domain Rate (TDR). TDR measures the percentage of transported samples that are classified as belonging to the target domain, with a higher TDR indicating more successful domain transfer. For MNIST, we train a classifier on the full MNIST training set and evaluate it on the corresponding test set, obtaining a baseline classification accuracy of 0.9905. For CelebA, we use the FairFace classifier [42], which is trained on a separate face dataset containing 108,501 images and does not use CelebA as its training data.

Target Probability Score (TPS). TPS measures the average probability(confidence) that the classifier assigns to the target-domain classes.

For the MNIST experiment, the set of target-domain labels

$$\mathcal{Y}_{\text{target}} = \{5, 6, 7, 8, 9\}.$$

For each transported sample $T(x_i)$, the classifier outputs the class probabilities

$$p_{\theta}(y = j \mid T(x_i)), \quad j \in \{0, \dots, 9\}.$$

The probability mass assigned to the target domain for the i -th transported sample is

$$q_i = \sum_{j \in \mathcal{Y}_{\text{target}}} p_{\theta}(y = j \mid T(x_i)).$$

TPS is then defined as

$$\text{TPS} = \frac{1}{N} \sum_{i=1}^N \sum_{j \in \mathcal{Y}_{\text{target}}} p_{\theta}(y = j \mid T(x_i)),$$

where N denotes the total number of transported samples. A higher TPS indicates that the classifier is more confident in assigning the transported samples to the target domain.

Fréchet Inception Distance (FID). FID [38] measures the discrepancy between the distribution of transported samples and that of real target-domain samples in a pretrained feature space. For MNIST, transported samples and real target samples are processed by the feature extractor of a pretrained MNIST classifier, and the distance is computed by:

$$\text{FID}_{\text{MNIST}} = \|\mu_{\text{map}} - \mu_{\text{target}}\|_2^2 + \text{Tr}\left(\Sigma_{\text{map}} + \Sigma_{\text{target}} - 2(\Sigma_{\text{map}}\Sigma_{\text{target}})^{1/2}\right).$$

where μ and Σ denote the empirical mean and covariance matrix, respectively, of the features extracted from the samples. For CelebA, we use `pytorch-fid` implementation [46].

Class-distribution Total Variation Distance (Class-TV). Class-TV measures how closely the class distribution of the generated target samples matches the class distribution of the real target dataset over digits 5-9. Let p_k denote the proportion of generated samples predicted as class k , conditioned on the prediction belonging to the target classes $\{5, \dots, 9\}$, and let q_k denote the proportion of class k in the real target test set. The metric is defined as

$$\text{TV}(p, q) = \frac{1}{2} \sum_{k=5}^9 |p_k - q_k|.$$

A smaller value indicates better agreement between the generated and real target-class distributions.

maximum mean discrepancy (MMD). MMD [39] measures the discrepancy between the transported distribution and the target distribution. It is computed by Gaussian RBF kernels. $\{x_i\}_{i=1}^n$ denote the target samples and $\{y_j\}_{j=1}^m$ denote the transported samples. For a kernel

$$k_\gamma(u, v) = \exp(-\gamma\|u - v\|_2^2),$$

the empirical squared MMD is

$$\widehat{\text{MMD}}_\gamma^2 = \frac{1}{n^2} \sum_{i, i'} k_\gamma(x_i, x_{i'}) + \frac{1}{m^2} \sum_{j, j'} k_\gamma(y_j, y_{j'}) - \frac{2}{nm} \sum_{i, j} k_\gamma(x_i, y_j).$$

We evaluate this quantity over 50 logarithmically spaced kernel parameters $\gamma \in [10^{-3}, 10]$ and report their average. A smaller MMD value indicates closer agreement between the transported and target distributions.

Enrichment-k100. Enrichment-k100 assesses the mixing between the transported and real target samples. For each transported sample x_i , we identify its k nearest neighbors $\mathcal{N}_k(x_i)$, in the combined transported and target sample set using Euclidean distance, excluding the query sample itself. The enrichment score is defined as

$$\text{Enrichment}_k = \frac{1}{nk} \sum_{i=1}^n \sum_{z \in \mathcal{N}_k(x_i)} \mathbf{1}\{z \text{ is transported}\}.$$

We use $k = 100$. When the transported and target sets have equal sizes, the random-mixing baseline is approximately 0.5. Therefore, values closer to 0.5 indicate better local mixing, whereas larger values indicate that transported samples cluster separately from the real target samples.

C.1.3 Efficiency

Training Time Training time is measured as the elapsed training time required to reach the checkpoint used for quantitative evaluation. It includes the training loop, data loading, optimizer updates, logging, and checkpointing, but excludes post-training evaluation such as FID, classifier-based metrics, and kNN enrichment. For fairness, all methods were trained using the same batch size and learning rate, with Adam optimizer and learning rate (2×10^{-4}). MNIST, CelebA, and Cell experiments were run on a single NVIDIA GeForce RTX 5060 Laptop GPU with 8 GB memory, using PyTorch 2.8.0+cu128, CUDA 12.8, and Python 3.9.25.

Female/Male Image Exposures. Female/Male Image Exposures count how many times real female/male images from the CelebA training set are used during training. These metrics provide a standardized measure of sample usage across methods, with fewer exposures indicating greater sample efficiency. In Neural OT, FIE and MIE differ because its multiple inner transport-map updates repeatedly access source female images without requiring additional target male images.

D Additional Experimental Results

D.1 MNIST

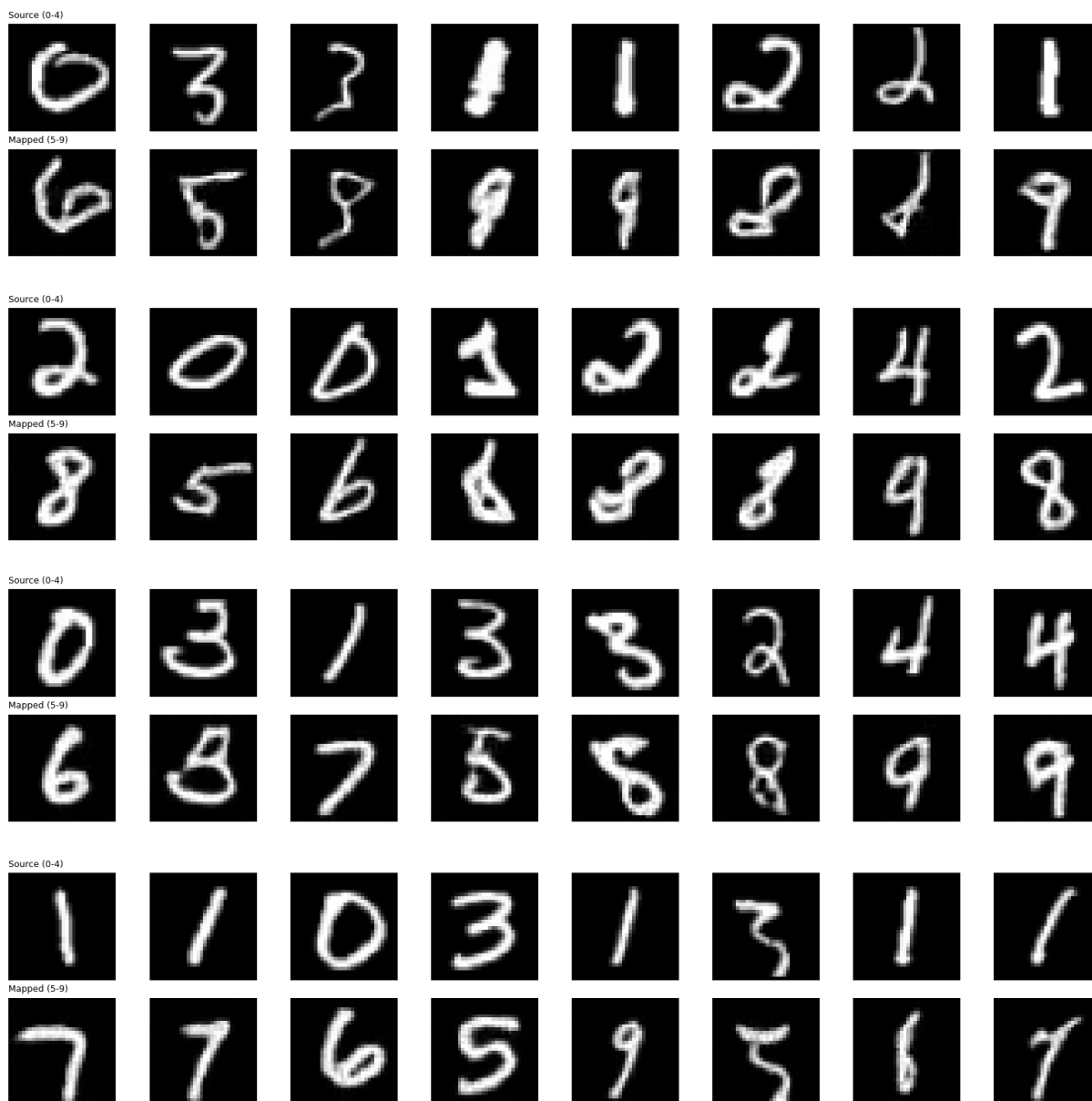


Figure 10: Additional qualitative results on the MNIST Digit Transfer (from digit 0-4 to digit 5-9) translation task.

D.2 CelebA



Figure 11: Qualitative results using the first 50 female images from the CelebA dataset as fixed source samples. Odd-numbered rows show the source female images, and the corresponding even-numbered rows show the mapped male images generated by CyclOT. The full image grid is displayed without post-hoc cropping.



Figure 12: Qualitative results using the first 50 non-smiling images from the CelebA dataset as fixed source samples. Odd-numbered rows show the source non-smiling images, and the corresponding even-numbered rows show the mapped smiling images generated by CyclOT. The full image grid is displayed without post-hoc cropping.

D.3 Pneumonia

Quantitative results in 6.6 only included the transport from healthy to unhealthy images. Table 7 includes results in both directions, where $Pos \rightarrow Neg$ indicates transporting from healthy to unhealthy, and $Neg \rightarrow Pos$ indicates transporting from unhealthy to healthy. Baseline predictions for true images in each class can be found in Table 6.

Method	Direction	Cost $\downarrow$	$ \text{Enrich-k100} - 0.5 \downarrow$	TPS $\uparrow$
ICNN OT	Pos $\rightarrow$ Neg	73142.1836	0.4986	0.1517 ± 0.0260
	Neg $\rightarrow$ Pos	65095.5671	0.4984	0.1513 ± 0.0409
Neural OT	Pos $\rightarrow$ Neg	1623.5243	0.0507	0.4180 ± 0.1811
	Neg $\rightarrow$ Pos	1984.1612	0.3246	0.4762 ± 0.1860
Ours	Pos $\rightarrow$ Neg	1472.9777	0.0133	0.4484 ± 0.1856
	Neg $\rightarrow$ Pos	1717.2522	0.3169	0.4098 ± 0.1873

Table 7: Quantitative comparison between ICNN OT, Neural OT, and our method on the chest x-ray healthy to pneumonia transition task, broken down by transport direction.

D.4 Ablation on cycle-consistency loss

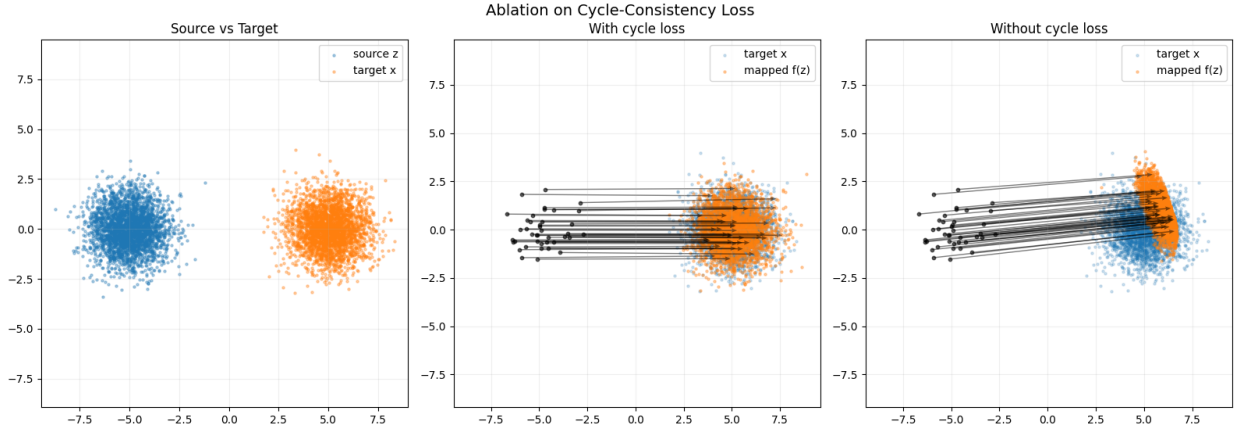


Figure 13: Ablation study on the cycle-consistency loss. Removing the cycle loss leads to mode collapse transport (right) compared to the full model (middle).