

Inverse Entropic Optimal Transport Solves Semi-supervised Learning via Data Likelihood Maximization

Mikhail Pershianov^{1,2} Arip Asadulaev³ Nikita Andreev Nikita Starodubcev⁴
Dmitry Baranchuk⁴ Anastasis Kratsios^{5,6} Evgeny Burnaev^{1,7} Alexander Korotin^{1,7}

¹Applied AI Institute

²Yandex School of Data Analysis

³MBZUAI

⁴Yandex Research

⁵Vector Institute

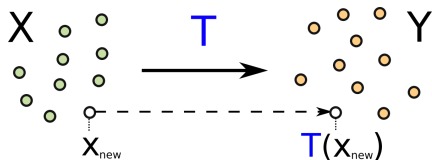
⁶McMaster University

⁷AXXX

Introduction to Semi-Supervised Domain Translation (SSDT)

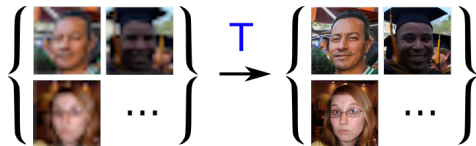
Motivation

The task: learn (from samples) a *translation* map between the two given data domains.

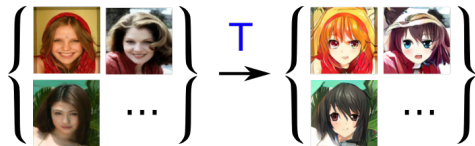


Important: the map should generalize to new data (similar to the train set).

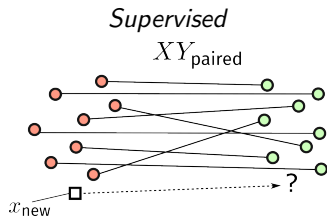
Example 1: Image Super-Resolution



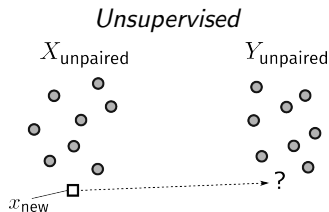
Example 2: Style Translation



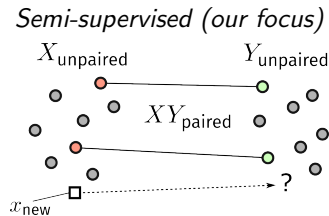
Problem Settings



- + Relatively straightforward to solve.
- Paired data collection is **expensive**.



- + Collecting unpaired data is straightforward.
- **Ambiguity** in solutions.



- + Combines strengths of both sides.
- No principled learning algorithms.

Applications of Semi-Supervised Domain Translation



Figure 1: Day $\rightarrow$ Night translation for autonomous driving³.



Figure 2: Sketch $\rightarrow$ Image translation for image editing³.

Other examples:

- **Medical segmentation:** Few labeled MRI/CT scans + unlabeled target scans¹.
- **Remote sensing:** Land-cover or cloud detection with few labeled satellite images².
- **Cross-modality:** LiDAR $\rightarrow$ RGB / Night $\rightarrow$ Day for perception or surveillance³.

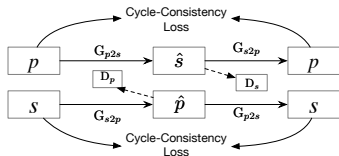
¹Xiaofeng Liu et al. (2022). “Act: Semi-supervised domain-adaptive medical image segmentation with asymmetric co-training”. In: *MICCAI*. Springer, pp. 66–76.

²Yadang Chen et al. (2022). “Semi-supervised contrastive learning for few-shot segmentation of remote sensing images”. In: *Remote Sensing* 14.17, p. 4254.

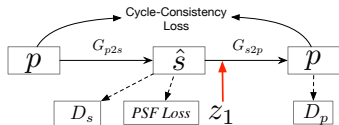
³Gaurav Parmar et al. (2024). “One-Step Image Translation with Text-to-Image Models”. In: *arXiv e-prints*.

Limitations of Existing Approaches

- Overcomplicated **heuristic** loss term construction $\rightarrow$ many hyperparameters.
- **No guarantees** for convergence to ground-truth solution.



(a) Unpaired CycleGAN architecture



(b) Semi-CycleGAN architecture

Figure 3: Complicated objective⁴.

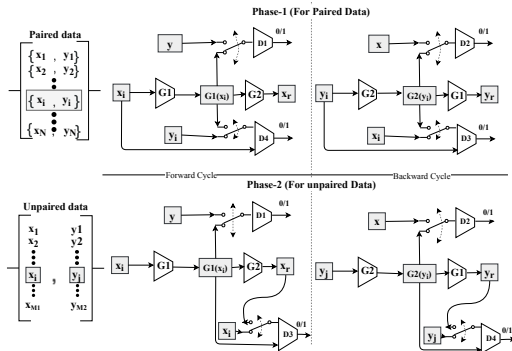


Figure 4: Complicated pipeline⁵.

⁴Chaofeng Chen et al. (2023). "Semi-supervised Cycle-GAN for face photo-sketch translation in the wild". In: *Computer Vision and Image Understanding* 235, p. 103775

⁵Soumya Tripathy, Juho Kannala, and Esa Rahtu (2019). "Learning image-to-image translation using paired and unpaired training samples". In: *Computer Vision-ACCV 2018*. Springer, pp. 51–66

Our Proposed Energy-Based Approach for SSDT

Learning Setup for SSDT

Idea: Semi-supervised domain translation leverages both paired and unpaired data to improve performance and generalization. Paired data provides precise guidance, while unpaired data helps the model adapt to abundant or unlabeled samples.

Assumption: We have access to:

- Paired samples: $XY_{\text{paired}} \sim \pi^*$
- Additional unpaired samples: $X_{\text{unpaired}} \sim \pi_x^*$, $Y_{\text{unpaired}} \sim \pi_y^*$

Setup: The goal is to learn the true conditional mapping $\pi^*(\cdot | x)$ using all available data.

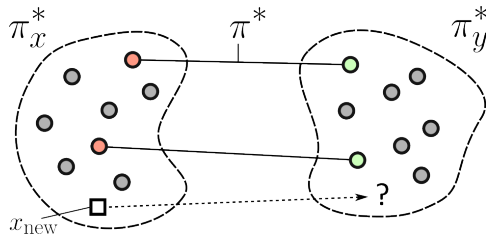


Figure 5: *Semi-supervised* domain translation setup (our focus).

Part I: Data Likelihood Maximization

Goal: Approximate the true distribution π^* by a parametric model π^θ via KL minimization:

$$\text{KL}(\pi^* \parallel \pi^\theta) = \underbrace{\text{KL}(\pi_x^* \parallel \pi_x^\theta)}_{\text{Marginal}} + \underbrace{\mathbb{E}_{x \sim \pi_x^*} \text{KL}(\pi^*(\cdot|x) \parallel \pi^\theta(\cdot|x))}_{\text{Conditional}}. \quad (1)$$

The conditional KL is crucial for domain translation since we often sample from marginal π_x^* and then sample from $\pi^\theta(\cdot|x)$. It can be rewritten as:

$$\mathbb{E}_{x \sim \pi_x^*} \mathbb{E}_{y \sim \pi^*(\cdot|x)} [\log \pi^*(y|x) - \log \pi^\theta(y|x)] = - \underbrace{\mathbb{E}_{x \sim \pi_x^*} H(\pi^*(\cdot|x))}_{=\text{const}} - \mathbb{E}_{x, y \sim \pi^*} \log \pi^\theta(y|x).$$

where the first term is independent of θ .

Resulting minimization objective equivalent to maximizing the conditional likelihood⁶:

$$\mathcal{L}(\theta) \stackrel{\text{def}}{=} -\mathbb{E}_{x, y \sim \pi^*} \log \pi^\theta(y|x) \rightarrow \min. \quad (2)$$

Limitation: Requires **fully paired data**; unpaired samples cannot be used.

In this paper, we **remove** this restriction by leveraging both paired and unpaired data.

⁶George Papamakarios et al. (2021). "Normalizing flows for probabilistic modeling and inference". In: *Journal of Machine Learning Research* 22.57, pp. 1–64.

Part II: Exploiting Unpaired Data via Energy-Based Parameterization

To address the above-mentioned issue and utilize unpaired data, we first use **Gibbs-Boltzmann parametrization**⁷:

$$\pi^\theta(y|x) \stackrel{\text{def}}{=} \frac{\exp(-E^\theta(y|x))}{Z^\theta(x)}, \quad Z^\theta(x) \stackrel{\text{def}}{=} \int_{\mathcal{Y}} \exp(-E^\theta(y|x)) \, dy, \quad (3)$$

where $E^\theta(\cdot|x) : \mathcal{Y} \rightarrow \mathbb{R}$ is the *Energy function*, and $Z^\theta(x)$ is the *normalization constant*. Substituting (3) into (2), we obtain:

$$\mathcal{L}(\theta) = \underbrace{\mathbb{E}_{x,y \sim \pi^*} E^\theta(y|x)}_{\text{Conditional, requires pairs } (x,y) \sim \pi^*} + \underbrace{\mathbb{E}_{x \sim \pi_x^*} \log Z^\theta(x)}_{\text{Marginal, requires } x \sim \pi_x^*} \rightarrow \min_{\theta}. \quad (4)$$

Observation: This objective already provides an opportunity to **exploit the unpaired samples** from the marginal distribution π_x^* to learn the conditional distributions $\pi^\theta(\cdot|x) \approx \pi^*(\cdot|x)$.

⁷Yann LeCun et al. (2006). “A tutorial on energy-based learning”. In: *Predicting structured data* 1.0.

Decoupling Cost and Potential in Energy Function

Goal: Incorporate independent samples $y \sim \pi_y^*$ into (4) to enable semi-supervised learning.

We propose, the following **energy function parameterization** for $\varepsilon > 0$:

$$E^\theta(y|x) \stackrel{\text{def}}{=} \frac{c^\theta(x, y) - f^\theta(y)}{\varepsilon}, \quad (5)$$

which **decouples** contributions from paired samples $(x, y) \sim \pi^*$ and unpaired samples $y \sim \pi_y^*$.

Substituting (5) into (4) and using $\mathbb{E}_{x, y \sim \pi^*} f^\theta(y) = \mathbb{E}_{y \sim \pi_y^*} f^\theta(y)$, we obtain our **final objective**:

$$\mathcal{L}(\theta) = \underbrace{\varepsilon^{-1} \mathbb{E}_{x, y \sim \pi^*} [c^\theta(x, y)]}_{\text{Joint, requires pairs } (x, y) \sim \pi^*} - \underbrace{\varepsilon^{-1} \mathbb{E}_{y \sim \pi_y^*} f^\theta(y)}_{\text{Marginal, requires } y \sim \pi_y^*} + \underbrace{\mathbb{E}_{x \sim \pi_x^*} \log Z^\theta(x)}_{\text{Marginal, requires } x \sim \pi_x^*} \rightarrow \min_{\theta}. \quad (6)$$

Question: How can we optimize this objective efficiently, given the typically **intractable** $Z^\theta(x)$?

Proposed Practical Parametrization: EBiEOT-GMM

Cost function and dual potential parameterization:

$$c^\theta(x, y) = -\varepsilon \log \sum_{m=1}^M v_m^\theta(x) \exp \left(\frac{\langle a_m^\theta(x), y \rangle}{\varepsilon} \right), \quad (7)$$

$$f^\theta(y) = \varepsilon \log \sum_{n=1}^N w_n^\theta \mathcal{N}(y \mid b_n^\theta, \varepsilon B_n^\theta), \quad (8)$$

where $v_m^\theta(x) : \mathbb{R}^{D_x} \rightarrow \mathbb{R}_+$ and $a_m^\theta(x) : \mathbb{R}^{D_x} \rightarrow \mathbb{R}^{D_y}$ are neural networks with learnable parameters θ_c and $\theta_f = \{w_n^\theta, b_n^\theta, B_n^\theta\}_{n=1}^N$ are learnable parameters.

Proposition (**Tractable** normalization constant)

Our parameterization of the cost c^θ and dual potential f^θ yields a closed-form $Z^\theta(x)$, enabling efficient end-to-end SGD⁸ optimization of the loss (6), given by:

$$Z^\theta(x) \stackrel{\text{def}}{=} \sum_{m=1}^M \sum_{n=1}^N z_{mn}(x), \text{ where } z_{mn}(x) \stackrel{\text{def}}{=} w_n v_m(x) \exp \left(\frac{a_m(x)^\top B_n a_m(x) + 2b_n^\top a_m(x)}{2\varepsilon} \right).$$

⁸Diederik P Kingma (2014). "Adam: A method for stochastic optimization". In: *ICLR, 2014*.

Training: Minimize the empirical counterpart of $\mathcal{L}(\theta)$ using SGD:

$$\mathcal{L}(\theta) \approx \hat{\mathcal{L}}(\theta) \stackrel{\text{def}}{=} \varepsilon^{-1} \frac{1}{P} \sum_{p=1}^P c^\theta(x_p, y_p) - \varepsilon^{-1} \frac{1}{R} \sum_{r=1}^R f^\theta(y_r) + \frac{1}{Q} \sum_{q=1}^Q \log Z^\theta(x_q).$$

Proposition (Tractable conditional distributions)

From our parametrization of the cost function (7) and dual potential (8) it follows that the $\pi^\theta(\cdot|x)$ are Gaussian mixtures:

$$\pi^\theta(y|x) = \frac{1}{Z^\theta(x)} \sum_{m=1}^M \sum_{n=1}^N z_{mn}(x) \mathcal{N}(y | d_{mn}(x), \varepsilon B_n), \quad (9)$$

where $d_{mn}(x) \stackrel{\text{def}}{=} b_n + B_n a_m(x)$ and $z_{mn}(x)$ is from Proposition 1.

Inference: Conditional distributions $\pi^\theta(\cdot|x)$ are Gaussian mixtures (9), enabling **efficient sampling** of y given x .

2D Translation: Gaussian $\rightarrow$ Swiss Roll. $P=128, Q=R=1024$

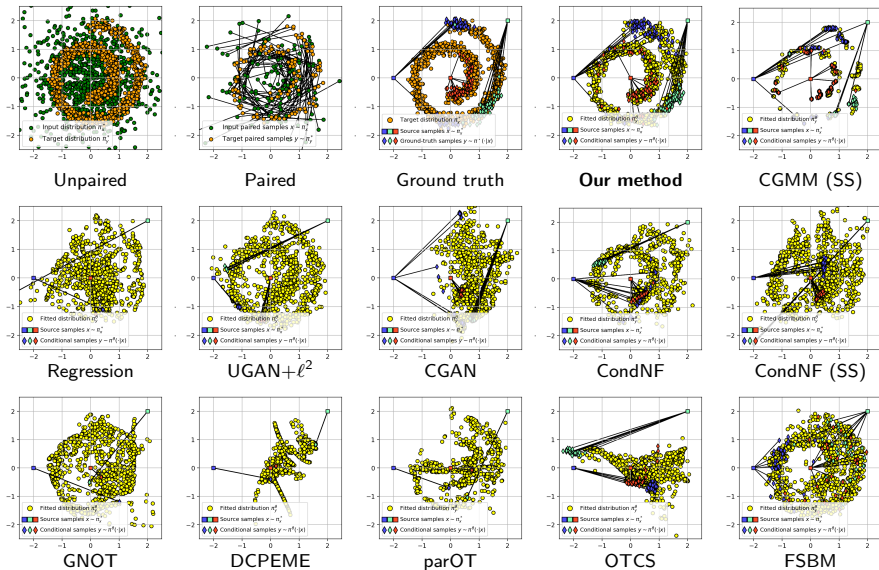


Image Translation via ALAE⁹(setup from FSBM paper¹⁰)

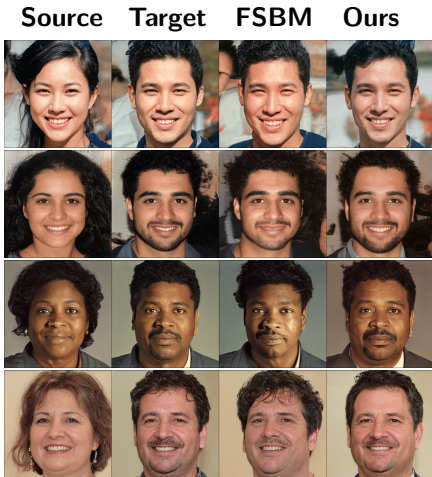


Figure 6: *Woman-to-Man* translation task.



Figure 7: *Old-to-Young* translation task.

⁹Panagiotis Theodoropoulos et al. (2025). “Feedback Schrödinger Bridge Matching”. In: *ICLR, 2025*.

¹⁰Stanislav Pidhorskyi, Donald A Adjeroh, and Gianfranco Doretto (2020). “Adversarial latent autoencoders”. In: *CVPR, 2020*.

Conclusion and Extensions

Discussion and Conclusion

Contribution:

- First ever **simple, non-minimax objective** integrating paired and unpaired data.
- Supports two parameterizations for continuous $y \in \mathbb{R}^{D_y}$: GMM and neural networks.

Future Directions:

- Extension to **discrete targets** $y \in \mathbb{K}^{D_y}$ is a promising direction.
- Established connection to iEOT may foster development of advanced semi-supervised methods based on **EBMs¹¹, diffusion schrödinger bridges, flow matching**.

Limitations:

- Relies on **GMM parameterization**, which affects scalability.
- Fully neural parameterizations for cost and potential relies on **MCMC sampling**.

Thank You!

Source code available at:



<https://github.com/MuXauJl111110/EBiEOT>

¹¹Yaxuan Zhu et al. (2024). “Learning Energy-Based Models by Cooperative Diffusion Recovery Likelihood”. In: *ICLR, 2024*.