

Explicit Density Approximation for Neural Implicit Samplers using a Bernstein-Based Convex Divergence

A Bernstein Ratio-Projection View of Likelihood-Free Density Learning from Rank Statistics

José Manuel de Frutos

Joint work with Pablo M. Olmos, Manuel A. Vázquez, Joaquín Míguez
Universidad Carlos III de Madrid

AISTATS 2026 Spotlight

Problem Setting: The "Explicitness" Gap

The Goal

Learn a sampler $\tilde{p}$ to mimic a target p given only finite samples $\{x_i\} \sim p$.

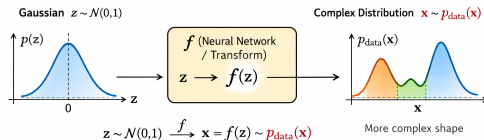
The Current Challenge:

- **Non-Convexity:** Training lacks global convergence.
- **No Density:** Standard generators (GANs) are "black boxes" with no explicit $p(x)$.

Our Objective:

- **Convex Training** in the density space.
- **Explicit Approximation:** Closed-form density.
- **Likelihood-free:** Train using **only samples** from p .

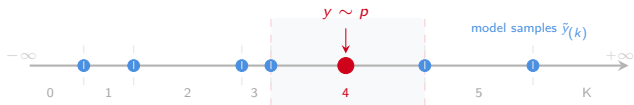
Entry Point: Rank statistics (Invariant Statistical loss ISL).



Can we recover the explicit density $\hat{p}(x)$ using only samples?

ISL: Rank Statistics to Uniform Histogram

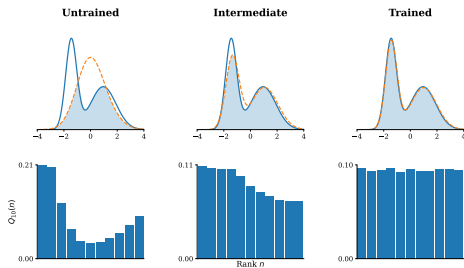
Is the data point y indistinguishable from the model samples $\{\tilde{y}_1, \dots, \tilde{y}_K\}$?



$$y \in (\tilde{y}_{(4)}, \tilde{y}_{(5)}] \implies \mathbf{A}_K(y) = 4$$



Aggregate N observations



Key Insight: If the model is perfect, the histogram of $A_K(y)$ is **uniform**.

ISL: From Rank Statistics to Discrepancy



The Rank Statistic

Count how many model points $\tilde{y}_j$ fall below the data point y :

$$A_K = \#\{j : \tilde{y}_j \leq y\}$$

The distribution of A_K is the **rank law**:

$$Q_K(n) = \mathbb{P}(A_K = n)$$

- ▶ If $p = \tilde{p}$, then A_K is **uniform**:

$$Q_K(n) = \frac{1}{K+1}$$

- ▶ Deviations indicate model mismatch.

ISL Discrepancy d_K

Measures the L^1 distance of the rank law to uniformity:

$$d_K(p, \tilde{p}) = \frac{1}{K+1} \sum_{n=0}^K \left| Q_K(n) - \frac{1}{K+1} \right|$$

Key Properties:

- ▶ **Convex** in its first argument p .
- ▶ **Weakly continuous** in p .
- ▶ Finite-rank approximation of an f -**divergence**.

Summary: $d_K = 0$ if and only if the rank histogram (Q_K) is exactly flat.

ISL Theory: From Ranks to Explicit Density

The Core Identity

Define the density ratio $q(x) = p(x)/\tilde{p}(x)$. The rank probabilities $Q_K(n)$ are actually the **Bernstein coefficients** of this ratio:

$$Q_K(n) = \int_0^1 b_{n,K}(t) q(F_{\tilde{p}}^{-1}(t)) dt$$

1. Bernstein Projection

Ranks provide the L^2 projection of the density ratio. Using the dual basis:

$$q_K(t) = \sum_{n=0}^K Q_K(n) \tilde{b}_{n,K}(t)$$

Result: Ranks $\Rightarrow$ projected ratio.

2. Density Estimator

We recover an explicit density $p_K(x)$ even from implicit models:

$$p_K(x) = \tilde{p}(x) \sum_{n=0}^K Q_K(n) \tilde{b}_{n,K}(F_{\tilde{p}}(x))$$

Result: $p_K(x) \rightarrow p(x)$ as $K \rightarrow \infty$.

Approximation Guarantee

$$\|q_K - 1\|_{\infty} \leq (K+1)^2 d_K(p, \tilde{p})$$

(includes quantitative convergence rates from Bernstein theory)

Multivariate Extension: Sliced Rank Discrepancy

The Sliced Mechanism

Project both distributions onto 1D for each direction $\xi \in \mathbb{S}^{d-1}$:

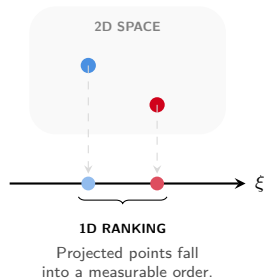
$$p_\xi = \xi_\# p, \quad \tilde{p}_\xi = \xi_\# \tilde{p}$$

Apply the 1D rank discrepancy and average over all directions:

$$d_K^{\text{slice}}(\tilde{p}, p) = \int_{\mathbb{S}^{d-1}} d_K(\tilde{p}_\xi, p_\xi) d\sigma(\xi)$$

Main Consequences:

- ▶ Preserves continuity and convexity.
- ▶ Practical in high dimensions.
- ▶ Natural fit for neural generators.

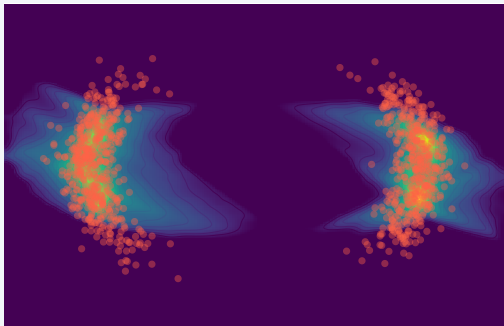


Interpretation

Scales the 1D rank idea to high-dimensional data via **random projections** Nadjahi et al. 2020.

Experiments: explicit density recovery and better mode coverage

2D: explicit density approximation



From the rank probabilities, we recover an explicit approximation p_K that tracks the target density well.

CelebA (64×64)

Method	FID↓	Prec.↑	Recall↑
DCGAN	30.93	0.839	0.834
dual-ISL + DCGAN	30.54	0.893	0.952
LS-DCGAN	22.99	0.932	0.524
dual-ISL + LS-DCGAN	23.56	0.928	0.824

Dual-ISL pretraining gives the strongest **mode coverage** while keeping fidelity competitive.

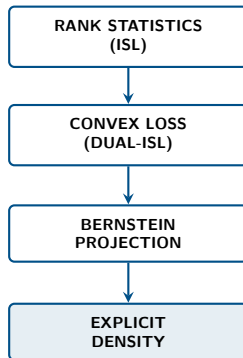
Empirical pattern

Dual-ISL gives

- ▶ **explicit density recovery** in 1D,
- ▶ **better recall / mode coverage**,
- ▶ and **smoother training** than standard ISL.

Take-home Message

- ▶ **The ISL Principle:** Uses ranks for likelihood-free distribution comparison.
- ▶ **Dual-ISL Swap:** Makes training strictly convex in the density.
- ▶ **Bernstein Identity:** Ranks are exactly the L^2 projections of the density ratio.
- ▶ **Density Recovery:** Enables explicit approximations for implicit models.
- ▶ **Scalability:** Slices scale to high-D and boost mode diversity.



Summary: Ranks $\rightarrow$ Convexity $\rightarrow$ Bernstein $\rightarrow$ Explicit Density.

Thank you for your attention!

Questions?

jofrutos@ing.uc3m.es



arXiv:2506.04700

References I

-  Frutos, José Manuel de, Pablo Olmos, et al. (2024). “Training implicit generative models via an invariant statistical loss”. In: *International Conference on Artificial Intelligence and Statistics*. PMLR, pp. 2026–2034.
-  Frutos, José Manuel de, Manuel A Vázquez, et al. (2024). “Robust training of implicit generative models for multivariate and heavy-tailed distributions with an invariant statistical loss”. In: *arXiv preprint arXiv:2410.22381*.
-  Nadjahi, Kimia et al. (2020). “Statistical and topological properties of sliced probability divergences”. In: *Advances in Neural Information Processing Systems* 33, pp. 20802–20812.
-  Stein, Viktor and José Manuel de Frutos (2026). “Approximating f -Divergences with Rank Statistics”. In: *arXiv preprint arXiv:2601.22784*.