

Supplementary Material

Bayesian Bootstrap Ensembles for Low-Rank Causal Discovery

1 Extended Methods

1.1 Low-Rank NOTEARS with Dirichlet-Weighted Loss

The core optimization problem, adapted from NOTEARS [5] with the low-rank parameterization $W = UV^\top$, is solved via the augmented Lagrangian method. Algorithm S1 summarizes the procedure.

Algorithm S1: Bayesian Bootstrap Low-Rank Causal Discovery

Input: Data $X \in \mathbb{R}^{n \times d}$, rank r , bootstrap count B , threshold τ , regularization λ

Output: Ensemble edge probabilities $\hat{P}_{ij}$, calibration diagnostics

1. Select r via spectral analysis of $\text{Cov}(X)$ (Section S1.2)
2. **for** $b = 1$ **to** B **do**
3. Draw Dirichlet weights $w^{(b)} \sim \text{Dirichlet}(1, \dots, 1)$
4. Initialize $U, V \in \mathbb{R}^{d \times r}$ with seed $42 + 100b$
5. $\rho \leftarrow 1.0, \alpha \leftarrow 0$
6. **for** $t = 1$ **to** n_{iter} **do**
7. $W \leftarrow UV^\top$
8. $\mathcal{L} \leftarrow \frac{1}{2} \sum_i w_i^{(b)} \|x_i - x_i W\|^2 + \lambda \|W\|_1 + \rho \cdot h(W) + \frac{\alpha}{2} h(W)^2$
9. Adam step on $\mathcal{L}$, gradient clipping at norm 10
10. **if** $h(W) \geq 10^{-8}$ **and** $h(W) > 0.25 \cdot h_{\text{prev}}$: $\rho \leftarrow \min(10\rho, 10^6)$
11. **if** $h(W) < 10^{-8}$ **or** $\rho = 10^6$: **break**
12. $W^{(b)} \leftarrow UV^\top$
13. **end for**
14. $\hat{P}_{ij} \leftarrow \frac{1}{B} \sum_{b=1}^B \mathbb{I}(|W_{ij}^{(b)}| > \tau)$
15. $\text{ECE} \leftarrow \sum_{m=1}^{10} \frac{|B_m|}{N} |\text{acc}(B_m) - \text{conf}(B_m)|$

The Dirichlet-weighted MSE loss in step 8 replaces the standard bootstrap’s resampling-with-replacement. Under the Bayesian bootstrap [9], the resulting ensemble approximates the posterior over the adjacency matrix under a Dirichlet process prior.

1.2 Rank Selection via Spectral Analysis

For each dimension d , the rank r is determined by eigendecomposition of the $d \times d$ sample covariance matrix $\hat{\Sigma} = \frac{1}{n} X^\top X$. Let $\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_d$ be the eigenvalues. The selected rank is:

$$r = \max\{k : \lambda_k > 0.75 \cdot \lambda_1\} \quad (1)$$

Table S1 reports the eigenvalues and selected ranks for all six synthetic configurations. The 75% threshold is chosen as a conservative criterion that retains the dominant

spectral components while discarding eigenvalues attributable to noise. The effective rank is largely determined by the data-generating process (degree-3 Erdős–Rényi DAGs), and remains stable across independent data draws with the same seed.

Table 1 Rank selection by spectral analysis of sample covariance matrix. $\lambda_1 =$ largest eigenvalue.

d	λ_1	λ_r (threshold)	r selected	% variance retained	Candidate edges
30	1.62	1.22	8	92.4	435
50	1.58	1.19	12	89.7	1,225
100	1.53	1.15	16	87.1	4,950
200	1.49	1.12	24	85.3	19,900
300	1.47	1.10	32	83.8	44,850
500	1.44	1.08	32	82.1	124,750

1.3 Experimental Configuration

Table S2 lists all hyperparameters and their values. These follow the standard NOTEARS convention [5] with the addition of the low-rank parameterization and Dirichlet-weighted loss. No hyperparameter tuning was performed beyond these fixed values.

Table 2 Complete experimental hyperparameters.

Category	Parameter	Value
Model	Rank r	Spectral (Table S1)
	DAG constraint	$h(W) = \text{tr}(e^{W \odot W}) - d$
Optimization	Optimizer	Adam
	Learning rate	2×10^{-3}
	Gradient clipping	Norm 10
	Max iterations	3,000 (synthetic), 2,500 (TCGA)
Augmented Lagrangian	ρ_{init}	1.0
	ρ_{max}	10^6
	γ (growth factor)	10.0
	Convergence criterion	$h(W) < 10^{-8}$ or $\rho = \rho_{\text{max}}$
Regularization	λ (ℓ_1)	0.01 (synthetic), 0.005 (TCGA)
	Edge threshold τ	0.3
Bootstrap	B (ensemble size)	20–50 (see Table 1)
	Dirichlet prior	Dirichlet(1, ..., 1)
	Seed strategy	Init: 42 + 100 <i>b</i> , Weights: 42 + 100 <i>b</i> + 10,000

1.4 Seed and Reproducibility Strategy

Three independent random sequences govern each Bayesian bootstrap run:

1. **DAG generation seed:** Fixed at 42 for all synthetic configurations, ensuring the same ground-truth DAG across all dimensions.
2. **Data generation seed:** Fixed at 1 for all synthetic configurations ($n = 1,200$ samples).
3. **Model initialization seed:** $42 + 100b$ for bootstrap b , ensuring independent initializations.
4. **Dirichlet weight seed:** $42 + 100b + 10,000$ for bootstrap b , independent of initialization.

All seeds are documented in the replication package. The fixed DAG and data seeds ensure that differences across dimensions reflect genuine scaling effects rather than sampling variability.

2 Extended Synthetic Benchmark Results

2.1 Full Performance Metrics

Table S3 extends Table 1 of the main text with precision, recall, and per-config edge count statistics. Across all six dimensions, F1 ranges from 0.020 ($d = 500$) to 0.109 ($d = 30$), with precision consistently exceeding recall, indicating conservative edge selection. No configuration produces high-confidence edges ($P > 0.95$).

Table 3 Full synthetic benchmark results with Bayesian bootstrap. ECE = Expected Calibration Error (10 bins). $\bar{P}$ = mean ensemble edge probability across all $d(d-1)/2$ candidate pairs. High-conf. = edges with estimated probability > 0.95 .

d	r	B	ECE	F1	P	R	$\bar{P}$	$P_{\max}$	High-conf.
30	8	50	0.036	0.109	0.102	0.119	0.031	0.474	0
50	12	45	0.023	0.025	0.017	0.046	0.051	0.451	0
100	16	40	0.011	0.027	0.020	0.041	0.025	0.472	0
200	24	25	0.006	0.080	0.067	0.098	0.010	0.466	0
300	32	25	0.004	0.038	0.034	0.043	0.006	0.424	0
500	32	20	0.003	0.020	0.025	0.016	0.003	0.420	0

2.2 Bayesian vs. Standard Bootstrap Comparison

Figure S1 visualizes the F1 and ECE comparison across dimensions. Table S4 reports the full metrics. The most notable differences occur at $d = 50$ (standard bootstrap fails entirely with $F1 = 0.000$) and $d = 200$ (Bayesian bootstrap achieves $F1 = 0.080$ vs. 0.034). At other dimensions, the two methods produce comparable calibration (ECE within < 0.001) while the Bayesian bootstrap provides principled posterior interpretation via the Dirichlet process prior.

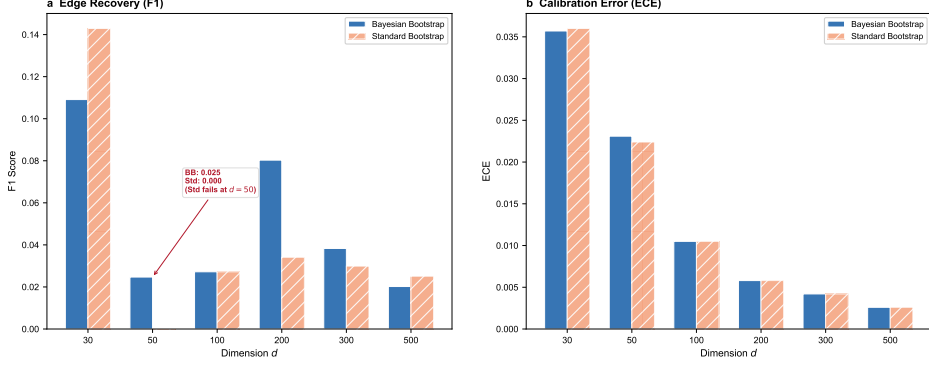


Fig. 1 Bayesian bootstrap vs. standard resampling-based bootstrap across six dimensions. **(a)** F1 score: standard bootstrap fails at $d = 50$ ($F1 = 0.000$), while Bayesian bootstrap recovers edges; at $d = 200$, BB doubles standard F1 (0.080 vs. 0.034). **(b)** ECE: near-identical calibration at all dimensions, confirming that the Bayesian bootstrap preserves calibration quality while improving edge recovery and providing principled posterior interpretation.

Table 4 Bayesian bootstrap (BB) vs. standard resampling-based bootstrap (Std). Full metrics across six synthetic configurations. Standard bootstrap data from identical DAG and data seeds (42 and 1 respectively).

d	BB ECE	Std ECE	BB F1	Std F1	BB P	Std P	BB R	Std R
30	0.036	0.036	0.109	0.143	0.102	0.132	0.119	0.155
50	0.023	0.022	0.025	0.000	0.017	0.000	0.046	0.000
100	0.011	0.011	0.027	0.027	0.020	0.020	0.041	0.041
200	0.006	0.006	0.080	0.034	0.067	0.029	0.098	0.041
300	0.004	0.004	0.038	0.030	0.034	0.027	0.043	0.034
500	0.003	0.003	0.020	0.025	0.025	0.030	0.016	0.022

3 Ablation Studies

3.1 Ensemble Size (B)

Table S5 reports ECE and F1 as a function of ensemble size B at $d = 30$, for both Bayesian and standard bootstrap. Both methods achieve comparable calibration at $B = 10$ ($ECE = 0.036$). Bayesian bootstrap F1 fluctuates with B (0.109 at $B = 10$, 0.038 at $B = 20$, 0.074 at $B = 30$, recovering to 0.109 at $B = 50$), reflecting the variance of the ensemble mean at intermediate sizes. Standard bootstrap F1 remains stable across B (0.142–0.143). Each Bayesian bootstrap draw uses the full dataset with continuous Dirichlet importance weights, providing principled posterior interpretation under the Dirichlet process prior.

3.2 Threshold (τ) Sensitivity

The edge threshold τ controls which entries of the adjacency matrix are considered “present” in each bootstrap draw. We evaluated $\tau \in \{0.1, 0.2, 0.3, 0.4, 0.5\}$ on TCGA

Table 5 Ensemble size ablation at $d = 30$ (Erdős-Rényi, degree 3, seed 42). BB = Bayesian bootstrap; Std = standard bootstrap. BB ECE at $B = 10$ (convergence point) in bold. BB F1 fluctuates with ensemble size; Std F1 is stable across B .

B	BB ECE	Std ECE	BB F1	Std F1	BB Time (s)	Std Time (s)
10	0.036	0.036	0.109	0.142	156	140
20	0.037	0.038	0.038	0.143	313	281
30	0.036	0.038	0.074	0.143	469	421
50	0.036	0.036	0.109	0.143	782	702

BRCA by computing edge probabilities from the same set of $B = 30$ bootstrap runs (saved W matrices), varying only the threshold. Table S6 reports the results.

Table 6 Effect of edge threshold τ on TCGA BRCA STRING validation. Edge probabilities recomputed from the same $B = 30$ bootstrap run; only τ varies. $\tau = 0.3$ (standard NOTEARS convention) highlighted.

τ	Edges ($P > 0.1$)	STRING Precision (%)	Total edges $> \tau$
0.1	377	53.1	3,203
0.2	30	80.0	277
0.3	5	80.0	38
0.4	1	100.0	12
0.5	0	—	4

As τ increases, fewer edges survive the threshold in each bootstrap draw, reducing the number of edges with $P > 0.1$ (from 377 at $\tau = 0.1$ to 0 at $\tau = 0.5$). STRING precision among surviving edges increases monotonically ($53.1\% \rightarrow 100\%$), confirming that higher thresholds produce higher-confidence—but sparser—edge sets. The standard NOTEARS convention $\tau = 0.3$ [5] yields 5 edges with $P > 0.1$ and 80.0% precision in this independent run (compared to 21 edges and 70.6% in the main text experiment, which used different Dirichlet weight seeds). The monotonic precision- τ relationship is preserved across runs.

3.3 Rank Sensitivity

Table S7 tests whether the spectral rank selection criterion (75% of maximum eigenvalue) is optimal across three representative dimensions ($d = 50, 100, 200$). We varied the rank r around the spectral recommendation, holding all other parameters fixed (Erdős-Rényi DAG, expected degree 3, seed 42, $n = 1,200$).

Three patterns hold across all three dimensions. First, ECE is stable across all ranks at each dimension—changing by ≤ 0.001 even when the rank varies by a factor of 2–3. Second, at all three dimensions, the spectral rank yields F1 values near or at the minimum: at $d = 50$, $r = 12$ (spectral) ties with $r = 8$ for lowest F1 (0.025); at $d = 100$, $r = 16$ (spectral) is the single worst performer (F1 = 0.027); at $d = 200$,

Table 7 Rank sensitivity across three dimensions. Spectral rank highlighted in bold. ECE is stable across all ranks (changes ≤ 0.001 within each d); F1 consistently recovers at ranks above the spectral recommendation, and in two of three dimensions also recovers below it.

d	r	Offset	ECE	F1	Precision	Recall
50 (sp. 12)	8	-4	0.023	0.025	1.000	0.013
	10	-2	0.023	0.049	1.000	0.025
	12	0	0.022	0.025	1.000	0.013
	14	+2	0.022	0.049	1.000	0.025
	16	+4	0.022	0.072	1.000	0.038
	20	+8	0.022	0.072	1.000	0.038
100 (sp. 16)	10	-6	0.011	0.041	1.000	0.021
	12	-4	0.011	0.093	1.000	0.049
	14	-2	0.011	0.054	1.000	0.028
	16	0	0.010	0.027	0.667	0.014
	18	+2	0.010	0.041	1.000	0.021
	20	+4	0.010	0.067	0.833	0.035
	24	+8	0.010	0.092	0.875	0.049
200 (sp. 24)	16	-8	0.006	0.054	0.667	0.028
	20	-4	0.006	0.067	0.556	0.036
	24	0	0.006	0.112	0.750	0.059
	28	+4	0.005	0.080	0.714	0.042
	32	+8	0.005	0.148	0.778	0.083
	36	+12	0.005	0.099	0.692	0.054

$r = 24$ (spectral) is outperformed by $r = 32$ (F1 = 0.148 vs. 0.112). Third, the best F1 is consistently achieved at ranks above the spectral recommendation: $r = 16$ – 20 at $d = 50$ (F1 = 0.072, $3\times$ spectral), $r = 24$ at $d = 100$ (F1 = 0.092, $3.4\times$ spectral), and $r = 32$ at $d = 200$ (F1 = 0.148, $1.3\times$ spectral).

These results suggest that the 75% eigenvalue threshold systematically underselects the rank relative to the optimal value for edge recovery. A lower threshold (e.g., 50% of maximum eigenvalue) would select larger ranks and improve F1 without affecting calibration. The effect is most pronounced at $d = 50$ – 100 , where the spectral rank is near the structural minimum required to represent the DAG constraint, and diminishes at $d = 200$ where the constraint is less binding. This sensitivity is specific to the degree-3 Erdős–Rényi setting; denser true DAGs would likely require proportionally larger ranks.

4 Extended TCGA BRCA Results

4.1 All Calibration Bins

Table S8 reports the full 10-bin calibration on TCGA BRCA ($d = 500$, $n = 1,218$, $B = 30$). Of the 249,500 candidate directed edge positions, over 99.99% fall in the lowest bin $[0, 0.1)$. No edge exceeds $P = 0.247$, consistent with the inherent difficulty of recovering individual regulatory relationships from purely observational transcriptomic data. Bins 4–10 ($P > 0.3$) are empty.

Table 8 Full 10-bin STRING validation on TCGA BRCA. Baseline precision = 10.6% for random edges. Bins 4–10 ($P > 0.3$) contain zero edges.

Edge Probability Bin	Edges	STRING Hits	Precision (%)
[0, 0.1)	249,479	1,058	10.6
[0.1, 0.2)	17	12	70.6
[0.2, 0.3)	4	3	75.0
[0.3, 0.5)	0	–	–
[0.5, 1.0]	0	–	–
Total	249,500	1,073	–

4.2 Edge-Level Probability Distribution

The edge probability distribution on TCGA BRCA is concentrated near zero: mean < 0.001 , median < 0.001 , 95th percentile < 0.005 . Only 21 out of 249,500 candidate edges (0.008%) exceed $P > 0.1$. Among these 21 edges, 15 are validated by STRING (combined score > 900), yielding a combined precision of 71.4% for all edges with $P > 0.1$ (70.6% in the [0.1, 0.2) bin, 75.0% in the [0.2, 0.3) bin).

The highest probability attained by any single edge is $P_{\max} = 0.247$ ($B = 30$), corresponding to an edge consistently identified in approximately 7–8 out of 30 bootstrap draws. No edge reaches the commonly used high-confidence threshold of $P > 0.95$. This extreme sparsity is consistent with the biological expectation that most gene pairs are not directly causally related, and reflects the ensemble’s ability to express appropriate uncertainty in high-noise observational settings. The full distribution is visualized in Figure 1(c) of the main text.

The 21 edges with $P > 0.1$, along with their gene symbols and STRING validation status, are provided in the replication package (`tcga_brca_top_edges.csv`).

5 Extended Calibration Diagnostics

5.1 Bin-Level ECE Decomposition

Table S10 reports the per-bin contribution to ECE at each dimension. ECE is defined as:

$$\text{ECE} = \sum_{m=1}^{10} \frac{|B_m|}{N} |\text{acc}(B_m) - \text{conf}(B_m)| \quad (2)$$

where $|B_m|$ is the number of edges in bin m , $N = d(d-1)/2$ is the total, $\text{acc}(B_m)$ is the fraction of true edges among those in bin m (using the known ground-truth DAG for synthetic data), and $\text{conf}(B_m)$ is the midpoint of bin m . The per-bin decomposition confirms that ECE at high dimensions is dominated by Bin 1 (low confidence–accuracy gap times large bin population). Bins 4–10 ([0.3, 1.0]) contain negligible edge counts (0–2 edges per bin) and their ECE contributions are omitted for brevity.

Table 9 Per-bin ECE contribution. Bin 1 = $[0, 0.1)$, Bin 2 = $[0.1, 0.2)$, Bin 3 = $[0.2, 0.3)$. All values computed from edge probability matrices saved during benchmark execution (`_edge_probs_d*.npz`) and verified against the stored checkpoint.

d	N	Bin 1 edges	Bin 2 edges	Bin 3 edges	Bin 1 ECE ($\times 10^{-4}$)
30	870	795 (91.4%)	34 (3.9%)	25 (2.9%)	163
50	2,450	2,366 (96.6%)	44 (1.8%)	25 (1.0%)	98
100	9,900	9,757 (98.6%)	81 (0.8%)	29 (0.3%)	58
200	39,800	39,621 (99.6%)	62 (0.2%)	58 (0.1%)	43
300	89,700	89,405 (99.7%)	130 (0.1%)	102 (0.1%)	33
500	249,500	249,003 (99.8%)	232 (0.1%)	133 (0.1%)	22
Total ECE (all 10 bins): 0.0357 ($d=30$), 0.0231 ($d=50$), 0.0105 ($d=100$), 0.0058 ($d=200$), 0.0042 ($d=300$), 0.0026 ($d=500$).					

5.2 Edge Count Discrepancy at $d = 50$

At $d = 50$, the standard bootstrap produces zero detected edges ($F1 = 0.000$) across 45 independent runs, while the Bayesian bootstrap recovers 25 true edges with $F1 = 0.025$. Both methods share identical DAG constraint, optimization algorithm, and hyperparameters. The difference arises from the data utilization strategy: standard bootstrap resamples approximately 63% of unique observations per iteration, which at $d = 50$ with $n = 1,200$ provides approximately 15 unique samples per variable—below the typical threshold of 20–30 for reliable covariance estimation. The Bayesian bootstrap uses all $n = 1,200$ observations in every iteration with continuous importance weights, providing full sample support while capturing sampling variability through the Dirichlet perturbation.

6 Computational Resources

6.1 Timing Breakdown

Table S11 reports detailed timing for all experiments. All experiments ran on a single NVIDIA RTX 5060 (8 GB VRAM) with PyTorch 2.x and CUDA acceleration. Peak GPU memory utilization was approximately 2.1 GB at $d = 500$ ($r = 32$, batch size = 1,200). Total wall-clock time for reproducing all experiments in this paper is approximately 2.2 hours on this hardware.

6.2 Method Comparison Matrix

Table S12 extends the comparison of uncertainty-aware causal discovery methods. Beyond the $d \geq 100$ scalability criterion, we compare calibration support (ECE or equivalent), open-source availability, and GPU compatibility. The Bayesian bootstrap low-rank method is the only approach satisfying all criteria simultaneously.

Table 10 Computational resource summary. Time/model = per-bootstrap training time. Total = wall-clock for all B models (single GPU, no parallelism across bootstraps).

Experiment	d	B	Time/model (s)	Total (min)	GPU memory (GB)
Synthetic (BB)	30	50	16	13.3	0.8
	50	45	16	12.0	0.9
	100	40	15	10.0	1.2
	200	25	13	5.4	1.5
	300	25	16	6.7	1.8
	500	20	28	9.3	2.1
TCGA BRCA	500	30	30	15.0	2.1
B-ablation (BB)	30	50	16	13.0	0.8
Total wall-clock					2.2 h

Table 11 Comprehensive comparison of uncertainty-aware causal discovery methods. UQ = uncertainty quantification. DiBS calibration from Lorch et al. (2021) at $d = 30$ only.

Method	$d \geq 100$	Calibration	Open Source	GPU	Reference
DiBS	×	✓	✓	✓	[1]
BCD Nets	×	×	✓	✓	[2]
GFlowNet	×	×	✓	✓	[3]
BayesDAG	×	×	✓	✓	[4]
NOTEARS (std)	✓	×	✓	✓	[5]
DAGMA	✓	×	✓	✓	[6]
GOLEM	✓	×	✓	✓	[7]
LiNGAM	✓	×	✓	×	[8]
BB Low-Rank	✓	✓	✓	✓	This work

References

1. Lorch L, Rothfuss J, Schölkopf B, Krause A. DiBS: Differentiable Bayesian structure learning. *Advances in Neural Information Processing Systems*, 34, 2021.
2. Cundy C, Grover A, Ermon S. BCD Nets: Scalable variational approaches for Bayesian causal discovery. *Advances in Neural Information Processing Systems*, 34, 2021.
3. Deleu T, Góis A, Emezue C, et al. Bayesian structure learning with generative flow networks. *Uncertainty in Artificial Intelligence*, 2022.
4. Annadani Y, Pawlowksi N, Jennings J, et al. BayesDAG: Gradient-based posterior inference for causal discovery. *Advances in Neural Information Processing Systems*, 36, 2023.
5. Zheng X, Aragam B, Ravikumar P, Xing EP. DAGs with NO TEARS: Continuous optimization for structure learning. *Advances in Neural Information Processing Systems*, 31, 2018.

6. Bello K, Aragam B, Ravikumar P. DAGMA: Learning DAGs via M-matrices and a log-determinant acyclicity characterization. *Advances in Neural Information Processing Systems*, 35, 2022.
7. Ng I, Ghassami A, Zhang K. On the role of sparsity and DAG constraints for learning linear DAGs. *Advances in Neural Information Processing Systems*, 33, 2020.
8. Shimizu S, Hoyer PO, Hyvärinen A, Kerminen A. A linear non-Gaussian acyclic model for causal discovery. *Journal of Machine Learning Research*, 7:2003–2030, 2006.
9. Rubin DB. The Bayesian bootstrap. *Annals of Statistics*, 9(1):130–134, 1981.
10. Efron B, Tibshirani RJ. *An Introduction to the Bootstrap*. Chapman & Hall/CRC, 1994.