

Sorting-Based Ordering Priors for Causal Structure Learning

Changxin Rong Yanze Gao Taiyu Ban Xiangyu Wang* Huanhuan Chen*

School of Computer Science and Technology, University of Science and Technology of China

{sa312, hchen}@ustc.edu.cn

Abstract—Large language models offer a promising source of structural knowledge for causal structure learning. Among different priors, ordering priors are especially robust and actionable because they restrict feasible edge directions through a sequential order over variables. However, in iterative learning, eliciting globally consistent ordering priors under a limited query budget remains challenging. Querying only local relations in the current graph produces fragmentary constraints that can propagate structural bias, whereas extending queries to non-adjacent pairs incurs quadratic cost. We propose Sorting-Based Ordering Prior for Causal Structure Learning. The method reformulates structural verification as a path-based sorting problem. It decomposes the graph into maximal causal paths and exploits causal transitivity to rectify causal chains efficiently. In this way, it builds a globally consistent variable order with logarithmic query complexity. The elicited priors are injected through a constraint scheme that penalizes order-violating directions, while still allowing data-driven correction. The framework applies to both score-based and gradient-based methods. Experiments show improved structure recovery and faster convergence, yielding a better tradeoff between query cost and prior coverage.

Index Terms—Causal structure learning, large language model, ordering prior.

I. INTRODUCTION

Causal structure learning aims to recover causal relations as directed acyclic graphs (DAGs) from observational data [1], [2]. While fundamental to numerous applications [3], [4], purely data-driven recovery is often unstable due to noise and limited sample sizes [2], [5]. To mitigate this, incorporating prior knowledge has become a standard strategy [6]. Specifically, ordering priors constrain the search space by imposing a sequential order over variables, which is often more robust than edge constraints [7]–[9]. Recently, Large Language Models have emerged as powerful tools for inferring these priors directly from variable metadata [10]–[13].

Existing LLM-guided methods typically follow an iterative interaction paradigm [12]–[14]. However, eliciting priors in this loop faces a severe tradeoff between information coverage and query cost. Restricting queries to local, adjacent relations is cost-effective but yields fragmentary constraints that miss

This research was supported in part by the National Nature Science Foundation of China (No. 62406302, 62137002, 62576327), in part by the Natural Science Foundation of Anhui province (No. 2408085QF195), in part by the USTC Research Funds of the Double First-Class Initiative (Grant No. YD9110002085), in part by the Research Funds of Centre for Leading Medicine and Advanced Technologies of IHM (Grant No. 2025IHM01030).

* Corresponding authors

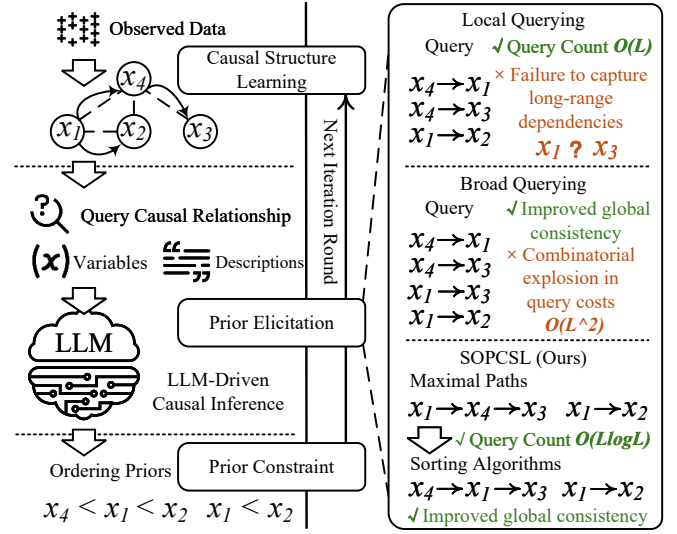


Fig. 1: Overview of the proposed SOPCSL framework. The method employs a sorting-based mechanism on maximal paths to identify ordering priors, securing a globally consistent sequential order with efficient log-linear complexity ($O(L \log L)$).

long-range dependencies, allowing structural biases to propagate. Conversely, expanding the query scope to non-adjacent pairs ensures global consistency but causes an unsustainable combinatorial explosion in query costs.

To resolve this dilemma, a fundamental shift from pairwise verification to sequence sorting offers a solution. Since causal DAGs are structurally composed of multiple causal chains representing the flow of influence, verifying a long path is equivalent to sorting its variables. Unlike pairwise checking, sorting algorithms inherently leverage causal transitivity¹. Consequently, a globally consistent sequential order can be constructed using only a logarithmic number of queries, effectively breaking the cost-coverage tradeoff.

Based on this insight, we propose Sorting-Based Ordering Prior for Causal Structure Learning (SOPCSL) (Fig. 1). At each iteration, SOPCSL decomposes the current graph into maximal causal paths and rectifies variable orders via sorting queries, utilizing majority voting to mitigate LLM hallucinations. The extracted ordering information is then translated into

¹If A causes B and B causes C, the relation A causes C is inferred automatically without query expenditure.

violation masks that penalize order-violating directions. This efficient constraint injection scheme seamlessly integrates into both score-based and gradient-based optimizations, providing high-quality global guidance with a minimal query budget.

The main contributions of this work are summarized as follows:

- We propose a sorting-based prior elicitation mechanism that transforms causal verification into a sequence sorting problem, leveraging transitivity to acquire globally consistent ordering constraints with high query efficiency.
- We develop a violation mask-based constraint injection scheme that integrates elicited priors into both score-based and gradient-based structure learning, balancing prior guidance with data-driven correction.
- Extensive experiments demonstrate that SOPCSL outperforms existing baselines in structure recovery accuracy and convergence speed under strict query budgets.

II. RELATED WORK

This section reviews prior-guided causal structure learning and the limitations of existing pairwise LLM-based causal discovery.

A. Causal Structure Learning with Prior Knowledge

Causal structure learning recovers Directed Acyclic Graphs (DAGs) from observational data via discrete score-based or continuous gradient-based optimization [1], [15]–[17]. However, purely data-driven recovery suffers from a super-exponential search space and limited finite-sample identifiability [18], [19]. Incorporating prior knowledge effectively mitigates these challenges [20]–[22]. Notably, ordering priors are highly valuable: knowing the topological order theoretically reduces the NP-hard structure learning problem to polynomial-time feature selection [18]. Thus, they provide a computationally efficient and theoretically grounded mechanism to guide modern causal discovery [23].

B. LLM-Driven Causal Discovery

LLMs are increasingly utilized as external knowledge bases for causal discovery [3], [24]. The predominant paradigm treats LLMs as pairwise reasoners, querying them to evaluate specific directed edges based on semantic metadata [24], [25]. Despite demonstrating the potential of semantic knowledge, this independent pairwise verification has intrinsic flaws. First, it ignores global consistency, potentially validating cyclical relations that violate DAG acyclicity assumptions. Second, constructing a globally consistent hierarchy without leveraging logical dependencies requires a computationally prohibitive quadratic number of queries. These limitations highlight the critical need for an efficient mechanism beyond pairwise checking to elicit globally consistent ordering priors—a gap our sorting-based framework aims to address.

III. PRELIMINARIES

This section introduces causal structure learning and the theoretical foundation of our method: DAGs, topological orderings, and maximal path decomposition.

A. Causal Structure Learning

Let $\mathcal{G} = (\mathbf{V}, \mathbf{E})$ be a directed acyclic graph (DAG) with d variables $\mathbf{V} = \{X_1, \dots, X_d\}$ and directed edges $\mathbf{E}$. Causal structure learning recovers $\mathcal{G}$ from observational data $\mathcal{D}$ via optimization under acyclicity constraints.

a) *Score-based Methods*: These methods maximize a scoring function σ (e.g., BIC or BDeu):

$$\max_{\mathcal{G}} \sum_{i=1}^d \sigma(X_i, pa(X_i), \mathcal{D}) \quad \text{s.t.} \quad \mathcal{G} \in \text{DAGs}, \quad (1)$$

where $pa(X_i)$ denotes the parents of X_i .

b) *Gradient-based Methods*: These methods minimize a reconstruction loss $\mathcal{L}$ by optimizing a weighted adjacency matrix $W \in \mathbb{R}^{d \times d}$, subject to a smooth acyclicity constraint $h(W) = 0$:

$$\min_W \mathcal{L}(W, \mathcal{D}) + \lambda \|W\|_1 \quad \text{s.t.} \quad h(W) = 0. \quad (2)$$

Both formulations face a super-exponential search space, making the incorporation of prior knowledge crucial for improving scalability and identifiability.

B. Topological Ordering as Prior

A topological ordering π of a DAG $\mathcal{G}$ is a linear sequence where every directed edge $X_i \rightarrow X_j$ satisfies $\pi(i) < \pi(j)$. A fixed ordering strictly restricts potential parents to predecessors, effectively reducing the search space from super-exponential to polynomial.

However, extracting globally consistent orderings from external knowledge (e.g., LLMs) is challenging. Independent local queries often introduce inaccuracies that propagate structural errors, necessitating an efficient mechanism to enforce global consistency.

C. Maximal Path Decomposition and Sorting

To efficiently verify graph structures, we decompose DAGs into maximal causal paths—sequences of distinct vertices $(v_1, \dots, v_k)$ where $(v_i, v_{i+1}) \in \mathbf{E}$.

Definition 1 (Maximal Path): A path $p = (v_1, \dots, v_k)$ in a DAG $\mathcal{G}$ is maximal if it cannot be extended at either end; meaning v_1 is a source node and v_k is a sink node.

Since a DAG is essentially supported by the union of its maximal paths, correctly identifying the causal order along these paths is sufficient to constrain the graph’s backbone.

This perspective transforms prior elicitation into a sorting problem. For a path of length L , exhaustive pairwise checks require $O(L^2)$ queries. In contrast, comparison-based sorting (e.g., Merge Sort) recovers a total order with only $O(L \log L)$ comparisons. This logarithmic efficiency motivates our strategy of sorting maximal paths to obtain globally consistent priors.

LLM Comparator Prompt Template

- 1. Task Definition:** You are an expert in causal inference. Your task is to determine the causal relations and topological ordering between two variables based on their definitions.
- 2. Variable Semantics:**
 - Variable A (s_i): [Definition and Unit]
 - Variable B (s_j): [Definition and Unit]
- 3. Query:** Considering the underlying mechanism, which variable naturally precedes the other in the causal hierarchy? Answer with “A precedes B” or “B precedes A” and provide a confidence score [0-1].

Fig. 2: Structured prompt template for the LLM comparator.

IV. METHODOLOGY

This section details SOPCSL, which reformulates prior elicitation as a topological sorting problem. By leveraging transitivity, it efficiently constructs consistent ordering priors and injects them into an iterative learning loop, aligning semantic knowledge with observational data for robust causal discovery.

A. Problem Setup

Let $V = \{X_1, \dots, X_d\}$ be d random variables and $\mathcal{D} = \{x^{(m)}\}_{m=1}^M$ be observational data. Causal structure learning aims to recover the true data-generating DAG $G^* = (V, E^*)$.

In our iterative paradigm, the learner produces an estimate G_t at iteration t using $\mathcal{D}$ and accumulated constraints. To correct spurious or reversed edges in G_t , prior knowledge is elicited from an LLM. Instead of computationally prohibitive broad querying, we leverage sorting-based transitivity to efficiently secure globally consistent priors.

B. Sorting-based Prior Elicitation

To efficiently refine G_t , we reformulate structural validation as a sequence sorting problem over maximal causal paths. Applying a sorting algorithm to these paths efficiently orders variables by their causal hierarchy, yielding globally consistent priors with log-linear query complexity ($O(L \log L)$).

Specifically, G_t is first decomposed into maximal causal paths $\mathbf{P}_t = \{p_1, \dots, p_k\}$. To recover the true hierarchy within each candidate chain p , a semantic sorting operator $\mathcal{S}$ (e.g., Merge Sort) reorders its nodes into a consistent permutation π_p :

$$\pi_p = \mathcal{S}(p \mid \mathcal{M}), \quad (3)$$

where $\mathcal{M}$ is an external semantic comparator invoked dynamically to compare variable pairs.

Instantiated by an LLM, the comparator $\mathcal{M}$ judges causal relations using variable semantics. Given a pair (X_i, X_j) with descriptions s_i and s_j , it outputs a directional prediction $y_{ij} \in$

$\{+1, -1\}$ (indicating $X_i \prec X_j$ or $X_j \prec X_i$) alongside a confidence score $r_{ij} \in [0, 1]$:

$$\mathcal{M}(X_i, X_j, s_i, s_j) \rightarrow (y_{ij}, r_{ij}). \quad (4)$$

The prompt template is shown in Fig. 2.

To update the global constraints, each sorted path $\pi_p = (v_1, v_2, \dots, v_m)$ is decomposed into adjacent directed pairs P_{π_p} :

$$P_{\pi_p} = \{(v_k, v_{k+1}) \mid 1 \leq k < m\}. \quad (5)$$

These pairs, retaining their confidence scores from the sorting process, are merged with the previous iteration’s constraints $\mathcal{O}_{t-1}$ to form a candidate pool $\mathcal{C}_t$:

$$\mathcal{C}_t = \mathcal{O}_{t-1} \cup \bigcup_{p \in \mathbf{P}_t} P_{\pi_p}. \quad (6)$$

Since merging independent paths may induce cycles, we employ a verification-based cycle-breaking strategy. Upon detecting a cycle in $\mathcal{C}_t$, we query each constituent edge multiple times to compute a robust aggregated reliability score S_e . We then strictly enforce acyclicity by removing the least reliable edge from the cycle:

$$\mathcal{C}_t \leftarrow \mathcal{C}_t \setminus \left\{ \arg \min_e S_e \right\}, \quad \text{until } \mathcal{C}_t \text{ is acyclic.} \quad (7)$$

The resulting acyclic set becomes the valid topological prior $\mathcal{O}_t$, ensuring mathematically consistent structural guidance grounded in reliable semantic judgments.

C. Prior Injection and Iterative Learning

To integrate the elicited ordering priors $\mathcal{O}_t$ into the learning loop, we encode them into a binary prior mask matrix $\mathbf{M}_O \in \{0, 1\}^{d \times d}$, where for any elicited relation $(X_i \prec X_j) \in \mathcal{O}_t$:

$$(\mathbf{M}_O)_{ij} = 1, \quad \text{and} \quad (\mathbf{M}_O)_{ji} = 0. \quad (8)$$

We adopt distinct constraint injection mechanisms for score-based and gradient-based backbones.

Injection for Score-based Learning. Score-based algorithms search over variable permutations $\pi \in \text{Perm}(V)$. We augment their standard scoring criterion with a penalty to discourage permutations that violate the partial order priors:

$$\min_{\pi \in \text{Perm}(V)} \left[\underbrace{\mathcal{F}(\pi; \mathcal{D})}_{\text{Structure Score}} + \beta \underbrace{\sum_{i,j} (\mathbf{M}_\pi)_{ij} \cdot (\mathbf{M}_O)_{ij}}_{\text{Prior Penalty } L_O} \right]. \quad (9)$$

Here, $\mathcal{F}(\pi; \mathcal{D}) = \min_{\mathcal{G} \in \mathcal{M}(\pi)} -\sigma(\mathcal{G}, \mathcal{D})$ is the optimal score achievable by any DAG consistent with order π . The penalty L_O quantifies prior conflicts using a violation mask $\mathbf{M}_\pi$, where $(\mathbf{M}_\pi)_{ij} = 1$ if $\text{pos}_\pi(i) > \text{pos}_\pi(j)$. A non-zero inner product with $\mathbf{M}_O$ penalizes arrangements that place X_i after X_j when the prior requires $X_i \prec X_j$.

Injection for Gradient-based Learning. Continuous optimization requires prohibiting all ancestral paths that contradict the prior, not just direct edges [23], [26]. To prevent indirect

TABLE I: Comparative performance analysis of structure recovery on score-based methods (BDeu and BIC). Performance is measured in terms of Scaled SHD and F1-score. The “Improvement” row indicates the percentage gain of SOPCSL relative to the Data-Only backbone.

Dataset	Cancer				Child				Insurance				Mildew			
Metrics	Scaled SHD↓		F1-score↑		Scaled SHD↓		F1-score↑		Scaled SHD↓		F1-score↑		Scaled SHD↓		F1-score↑	
Sample Size	250	1000	250	1000	500	2000	500	2000	500	2000	500	2000	8000	32000	8000	32000
BDeu	0.73	0.48	0.55	0.61	0.40	0.22	0.70	0.81	0.48	0.31	0.68	0.79	0.52	0.48	0.61	0.65
+Local Querying	0.69	0.44	0.58	0.64	0.38	0.20	0.72	0.84	0.47	0.30	0.68	0.80	0.50	0.46	0.63	0.67
+Broad Querying	0.57	0.32	0.69	0.76	0.24	0.06	0.87	0.98	0.42	0.25	0.72	0.86	0.43	0.13	0.74	0.93
+SOPCSL (Ours)	0.54	0.32	0.73	0.76	0.22	0.07	0.89	0.97	0.40	0.28	0.74	0.84	0.40	0.15	0.74	0.92
Improvement	26.0%	33.3%	32.7%	24.6%	45.0%	68.2%	27.1%	19.8%	16.7%	9.7%	8.8%	6.3%	23.1%	68.8%	21.3%	41.5%
BIC	0.98	0.65	0.42	0.49	0.36	0.19	0.73	0.82	0.65	0.58	0.52	0.55	0.76	0.75	0.41	0.48
+Local Querying	0.96	0.62	0.43	0.51	0.35	0.18	0.75	0.83	0.63	0.56	0.53	0.58	0.75	0.68	0.42	0.53
+Broad Querying	0.90	0.41	0.46	0.71	0.24	0.03	0.86	0.99	0.53	0.30	0.65	0.82	0.69	0.63	0.49	0.58
+SOPCSL (Ours)	0.93	0.37	0.44	0.75	0.27	0.03	0.83	0.99	0.50	0.33	0.67	0.79	0.72	0.65	0.46	0.56
Improvement	5.1%	43.1%	4.8%	53.1%	25.0%	84.2%	13.7%	20.7%	23.1%	43.1%	28.8%	43.6%	5.3%	13.3%	12.2%	16.7%

TABLE II: Comparative performance analysis of structure recovery on gradient-based frameworks. The “Improvement” row indicates the percentage gain of SOPCSL relative to the Data-Only backbone. Best results are highlighted in **bold**.

Dataset	Cancer				Child				Insurance				Mildew			
Metrics	Scaled SHD↓		F1-score↑		Scaled SHD↓		F1-score↑		Scaled SHD↓		F1-score↑		Scaled SHD↓		F1-score↑	
Sample Size	1d	10d	1d	10d	1d	10d	1d	10d	1d	10d	1d	10d	1d	10d	1d	10d
NOTEARS	0.92	0.28	0.45	0.81	0.88	0.35	0.49	0.78	0.95	0.82	0.38	0.46	0.85	0.62	0.48	0.65
+Local Querying	0.85	0.21	0.58	0.89	0.82	0.30	0.55	0.81	0.88	0.75	0.45	0.52	0.80	0.58	0.55	0.68
+Broad Querying	0.79	0.12	0.68	0.95	0.65	0.18	0.70	0.92	0.68	0.35	0.65	0.82	0.35	0.15	0.82	0.95
+SOPCSL (Ours)	0.78	0.10	0.67	0.96	0.68	0.21	0.68	0.90	0.65	0.37	0.68	0.80	0.37	0.18	0.80	0.92
Improvement	15.2%	64.3%	48.9%	18.5%	22.7%	40.0%	38.8%	15.4%	31.6%	54.9%	78.9%	73.9%	56.5%	71.0%	66.7%	41.5%
NOTEARS-MLP	1.60	0.85	0.28	0.55	1.55	0.55	0.22	0.65	1.35	0.62	0.15	0.58	1.45	0.58	0.15	0.62
+Local Querying	1.45	0.75	0.38	0.65	1.52	0.52	0.28	0.68	1.30	0.58	0.20	0.62	1.40	0.55	0.19	0.65
+Broad Querying	1.20	0.60	0.52	0.72	1.42	0.46	0.38	0.74	1.15	0.40	0.35	0.78	1.20	0.42	0.26	0.75
+SOPCSL (Ours)	1.25	0.58	0.49	0.75	1.40	0.48	0.35	0.75	1.12	0.42	0.32	0.79	1.22	0.45	0.24	0.72
Improvement	21.9%	31.8%	75.0%	36.4%	9.7%	12.7%	59.1%	15.4%	17.0%	32.3%	113%	36.2%	15.9%	22.4%	60.0%	16.1%

feedback loops (e.g., $X_j \rightarrow X_k \rightarrow X_i$ violating $X_i \prec X_j$), we regularize the graph’s transitive closure:

$$W_t = \arg \min_W [\mathcal{L}(W, \mathcal{D}) + \lambda \|W\|_1 + \gamma \mathcal{L}_O(W, \mathbf{M}_O)] \quad (10)$$

s.t. $h(W) = 0$,

where $\mathcal{L}$ is the reconstruction loss and $h(W) = 0$ ensures acyclicity. The prior penalty $\mathcal{L}_O(W, \mathbf{M}_O) = \|\mathbf{R}(W) \circ \mathbf{M}_O^T\|_1$ restricts the path reachability matrix $\mathbf{R}(W) = \sum_{k=1}^d |W|^k$. Penalizing its Hadamard product ($\circ$) with the transposed prior mask $\mathbf{M}_O^T$ forces the probability of any reverse path (flowing from X_j back to X_i) toward zero.

Iterative Learning Workflow. SOPCSL operates as a closed-loop active system initialized with $\mathcal{O}_0 = \emptyset$. At each iteration t :

- 1) **Structure Learning:** The learner solves Eq. (9) or Eq. (10) constrained by $\mathbf{M}_O$ to estimate G_t .
- 2) **Constraint Update:** Maximal paths extracted from G_t are sorted by the LLM and merged into $\mathcal{O}_{t+1}$ via the cycle-breaking mechanism.
- 3) **Termination:** The loop terminates when constraints stabilize ($\mathcal{O}_{t+1} = \mathcal{O}_t$).

V. EXPERIMENTS

This section empirically evaluates SOPCSL across four benchmark datasets, comparing it against data-only learning and alternative querying strategies. We also conduct ablations to analyze key components, LLM generalization, convergence speed, and query costs.

A. Datasets and Experimental Setup

We evaluate SOPCSL on four standard benchmark causal DAGs: Cancer, Child, Insurance, and Mildew. To assess robustness, we generate synthetic data under both small and large sample sizes, aligning the generation process with the respective learning backbones. For score-based methods, discrete data are sampled from Bayesian Networks; for gradient-based methods, continuous data are simulated using linear and non-linear SEMs with Gaussian noise. SOPCSL is integrated with diverse backbones: Tabu Search [20] (using BIC [27] and BDeu [28] scores) for discrete optimization, alongside NOTEARS [17], and the non-linear NOTEARS-MLP [29] for continuous optimization.

We primarily utilize GPT-4 [30] as the reasoning engine. To evaluate model generalization, a comparative analysis including Claude 4, Gemini 2.5 [31], DeepSeek V3.1 [32], and LLaMA 3.3 [33] is presented in Section V-D. Performance

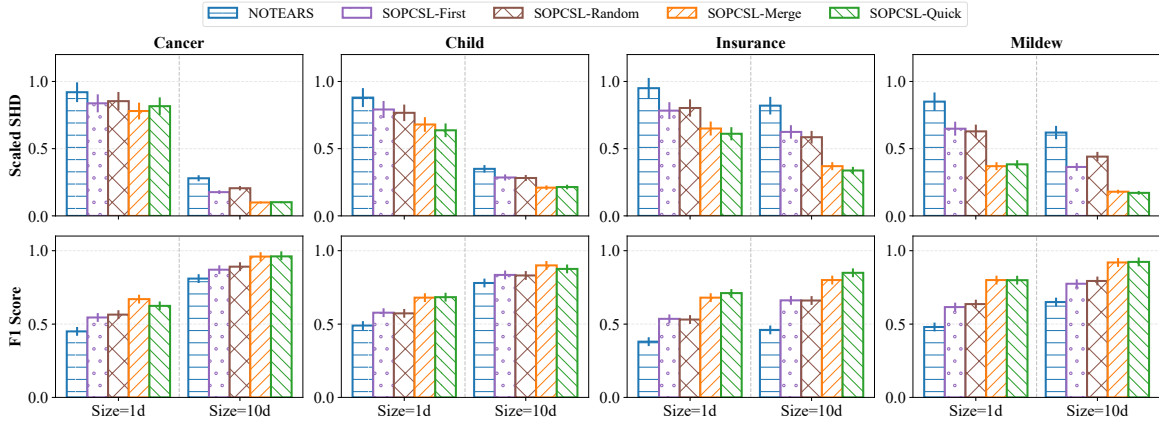


Fig. 3: Comparison of our method and other ablation methods in terms of F1-score and Scaled SHD.

TABLE III: Performance comparison of the proposed framework instantiated with different LLMs using NOTEARS as the structure learning backbone. The sample sizes are denoted as 1d and 10d. Best results are highlighted in **bold**.

Dataset	Cancer				Child				Insurance				Mildew			
	Scaled SHD↓		F1-score↑		Scaled SHD↓		F1-score↑		Scaled SHD↓		F1-score↑		Scaled SHD↓		F1-score↑	
	1d	10d	1d	10d	1d	10d	1d	10d	1d	10d	1d	10d	1d	10d	1d	10d
NOTEARS	0.92	0.28	0.45	0.81	0.88	0.35	0.49	0.78	0.95	0.82	0.38	0.46	0.85	0.62	0.48	0.65
+ Claude 4	0.83	0.18	0.61	0.90	0.76	0.25	0.66	0.87	0.83	0.55	0.53	0.70	0.55	0.33	0.69	0.85
+ DeepSeek V3.1	0.85	0.23	0.55	0.88	0.84	0.29	0.58	0.83	0.80	0.62	0.56	0.64	0.68	0.45	0.61	0.77
+ Gemini 2.5	0.79	0.05	0.64	0.98	0.72	0.23	0.65	0.89	0.63	0.45	0.70	0.77	0.42	0.22	0.78	0.90
+ GPT-4o	0.78	0.10	0.67	0.96	0.68	0.21	0.68	0.90	0.65	0.37	0.68	0.80	0.37	0.18	0.80	0.92
+ LLaMA 3.3	0.88	0.26	0.52	0.83	0.82	0.33	0.53	0.79	0.91	0.76	0.41	0.50	0.79	0.54	0.53	0.70

is quantified using the F1 score and the Scaled Structural Hamming Distance ($\text{SHD}/|E|$), which normalizes the error by the number of ground truth edges to facilitate fair comparisons across graphs of varying sizes.

B. Comparative Performance Analysis

We evaluate SOPCSL against a Data-Only baseline and two exhaustive querying paradigms: Local Querying (restricted to current graph edges) and Broad Querying (comprehensive pairwise checks). To isolate the strategy effect, all LLM-based methods utilize the same majority-voting conflict resolution and are run to convergence.

As shown in Tables I and II, SOPCSL achieves structure recovery accuracy comparable to the Broad Querying baseline. While Broad Querying outperforms the Local approach by expanding the search space, SOPCSL attains similar global consistency through sorting-based transitivity. This mechanism efficiently captures long-range dependencies, demonstrating that strategic path ordering provides structural guidance on par with brute-force enumeration.

C. Ablation Studies

We conduct ablations to examine algorithmic robustness and the necessity of conflict resolution. We compare the default Merge Sort with a Quick Sort variant, and benchmark the consensus-based cycle-breaking against First-Edge Pruning (deterministic removal) and Random-Edge Pruning (stochastic removal). As shown in Fig. 3, SOPCSL maintains strong performance across sorting algorithms, indicating low

sensitivity to the specific ordering protocol. The conflict resolution mechanism is critical. Unlike naive baselines that break contradictions arbitrarily and degrade accuracy, our consensus strategy uses repeated querying to prune low-confidence relations, preserving the coherence of the elicited prior.

D. Influence of Different LLMs

To evaluate generalization across reasoning engines, we instantiate SOPCSL with GPT-4o, Claude 4, Gemini 2.5, DeepSeek V3.1, and LLaMA 3.3, keeping all other settings fixed. As shown in Table III, while stronger models generally yield better recovery, the framework remains effective across all tested LLMs. This robustness stems from the sorting-based mechanism, which enforces transitivity to build consistent global orders, mitigating local reasoning errors in weaker models. Additionally, the soft constraint formulation allows data-driven evidence to override occasional hallucinations. Thus, SOPCSL successfully extracts valuable structural guidance even from models with moderate reasoning capabilities.

E. Iterative Dynamics and Efficiency Analysis

To characterize learning dynamics and resource consumption, we analyze the trajectories of SHD, F1 score, and single-round query counts across four benchmark datasets over successive learning iterations. This multi-dimensional evaluation benchmarks the proposed SOPCSL against Local Querying and Broad Querying baselines by plotting metric progression against interaction rounds to isolate the convergence speed and cost-effectiveness of each elicitation strategy.

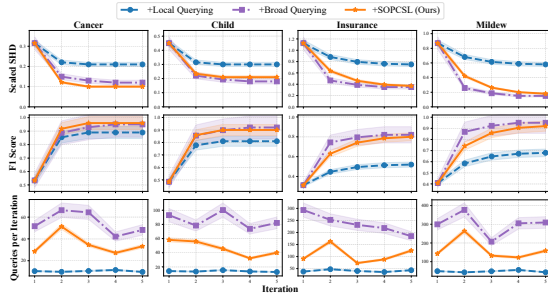


Fig. 4: Iterative dynamics of structure learning performance and elicitation cost across four datasets.

The results presented in Fig. 4 demonstrate that SOPCSL achieves an optimal convergence profile. In terms of structure recovery, the method exhibits rapid improvement and attains accuracy plateaus comparable to the Broad Querying baseline. Crucially, this high-fidelity recovery is achieved with a query growth rate similar to the Local baseline, standing in sharp contrast to the rapid cost escalation observed in the Broad baseline. This validates that the sorting-based mechanism effectively leverages transitivity to extract high-value constraints, thereby attaining the structural guidance of exhaustive search with the budgetary economy of local heuristics.

VI. CONCLUSION

This paper presented a novel sorting-based framework that reformulates prior elicitation as a topological sorting problem, constructing globally consistent ordering priors with log-linear query complexity. The approach ensures robustness, yielding structure recovery accuracy comparable to exhaustive querying protocols as demonstrated by comprehensive experiments. These findings confirm that strategic path ordering provides a scalable, high-fidelity mechanism for integrating large language model knowledge into causal discovery, effectively balancing computational efficiency with structural reliability.

REFERENCES

- [1] P. Spirtes, C. N. Glymour, R. Scheines, Causation, prediction, and search, MIT Press, 2000.
- [2] J. Peters, D. Janzing, B. Schölkopf, Elements of causal inference: foundations and learning algorithms, MIT Press, 2017.
- [3] E. Kıcıman, R. Ness, A. Sharma, C. Tan, Causal reasoning and large language models: Opening a new frontier for causality, arXiv preprint arXiv:2305.00050 (2023).
- [4] J. Ma, Causal inference with large language model: A survey, Findings of the Association for Computational Linguistics: NAACL 2025 (2025) 5886–5898.
- [5] M. Kalisch, P. Bühlman, Estimating high-dimensional directed acyclic graphs with the pc-algorithm., Journal of Machine Learning Research 8 (3) (2007).
- [6] A. C. Constantinou, Z. Guo, N. K. Kitson, The impact of prior knowledge on causal structure learning, Knowledge and Information Systems 65 (2023) 3385–3434.
- [7] M. Teyssier, D. Koller, Ordering-based search: A simple and effective algorithm for learning Bayesian networks, in: Proceedings of the 21st Conference on Uncertainty in Artificial Intelligence, 2005, pp. 584–590.
- [8] Y. W. Park, D. Klabjan, Bayesian network learning via topological order, Journal of Machine Learning Research 18 (99) (2017) 1–32.
- [9] T. Niinimäki, P. Parviainen, M. Koivisto, Structure discovery in bayesian networks by sampling partial orders, Journal of Machine Learning Research 17 (57) (2016) 1–47.

- [10] T. Ban, L. Chen, D. Lyu, X. Wang, H. Chen, Causal structure learning supervised by large language model, arXiv preprint arXiv:2311.11689 (2023).
- [11] L. Chen, T. Ban, X. Wang, D. Lyu, H. Chen, Mitigating prior errors in causal structure learning: Towards LLM driven prior knowledge, arXiv preprint arXiv:2306.07032 (2023).
- [12] T. Ban, L. Chen, D. Lyu, X. Wang, Q. Zhu, Q. Tu, H. Chen, Integrating large language model for improved causal discovery, IEEE Transactions on Artificial Intelligence (2025).
- [13] T. Ban, L. Chen, D. Lyu, X. Wang, Q. Zhu, H. Chen, Llm-driven causal discovery via harmonized prior, IEEE Transactions on Knowledge and Data Engineering (2025).
- [14] L. Chen, T. Ban, X. Wang, D. Lyu, H. Chen, Mitigating prior errors in causal structure learning: A resilient approach via bayesian networks, IEEE Transactions on Pattern Analysis and Machine Intelligence (2025).
- [15] J. Pearl, Causality: Models, Reasoning, and Inference, 2nd Edition, Cambridge University Press, Cambridge, UK, 2009.
- [16] G. F. Cooper, E. Herskovits, A Bayesian method for the induction of probabilistic networks from data, Machine Learning 9 (4) (1992) 309–347.
- [17] X. Zheng, B. Aragam, P. K. Ravikumar, E. P. Xing, DAGs with no tears: Continuous optimization for structure learning, in: Advances in Neural Information Processing Systems, Vol. 31, 2018, pp. 9492–9503.
- [18] D. M. Chickering, Learning Bayesian networks is NP-complete, Learning from data: Artificial intelligence and statistics V (1996) 121–130.
- [19] T. S. Verma, J. Pearl, On the equivalence of causal models, arXiv preprint arXiv:1304.1108 (2013).
- [20] A. Li, P. Beek, Bayesian network structure learning with side constraints, in: Proceedings of the International Conference on Probabilistic Graphical Models, 2018, pp. 225–236.
- [21] H. Amirkhani, M. Rahmati, P. J. Lucas, A. Hommersom, Exploiting experts’ knowledge for structure learning of Bayesian networks, IEEE Transactions on Pattern Analysis and Machine Intelligence 39 (11) (2016) 2154–2170.
- [22] X. Wang, T. Ban, L. Chen, D. Lyu, Q. Zhu, H. Chen, Large-scale hierarchical causal discovery via weak prior knowledge, IEEE Transactions on Knowledge and Data Engineering (2025).
- [23] T. Ban, L. Chen, X. Wang, X. Wang, D. Lyu, H. Chen, Differentiable structure learning with partial orders, in: The Thirty-eighth Annual Conference on Neural Information Processing Systems, 2024.
- [24] X. Liu, P. Xu, J. Wu, J. Yuan, Y. Yang, Y. Zhou, F. Liu, T. Guan, H. Wang, T. Yu, et al., Large language models and causal inference in collaboration: A comprehensive survey, arXiv preprint arXiv:2403.09606 (2024).
- [25] S. Long, A. Piché, V. Zantedeschi, T. Schuster, A. Drouin, Causal discovery with language models as imperfect experts, in: Proceedings of the ICML 2023 Workshop on Structured Probabilistic Inference & Generative Modeling, 2023.
- [26] L. Chen, T. Ban, D. Lyu, Y. Sun, K. Hu, X. Wang, H. Chen, Continuous structure constraint integration for robust causal discovery, in: The 28th International Conference on Artificial Intelligence and Statistics, 2025.
- [27] G. Schwarz, Estimating the dimension of a model, The Annals of Statistics 6 (1978) 461–464.
- [28] W. Buntine, Theory refinement on Bayesian networks, in: Uncertainty Proceedings, 1991, pp. 52–60.
- [29] X. Zheng, C. Dan, B. Aragam, P. Ravikumar, E. Xing, Learning sparse nonparametric dags, in: International Conference on Artificial Intelligence and Statistics, Pmlr, 2020, pp. 3414–3425.
- [30] J. Achiam, S. Adler, S. Agarwal, L. Ahmad, I. Akkaya, F. L. Aleman, D. Almeida, J. Altenschmidt, S. Altman, S. Anadkat, et al., GPT-4 technical report, arXiv preprint arXiv:2303.08774 (2023).
- [31] G. Comanici, E. Bieber, M. Schaeckermann, I. Pasupat, N. Sachdeva, I. Dhillon, M. Blistein, O. Ram, D. Zhang, E. Rosen, et al., Gemini 2.5: Pushing the frontier with advanced reasoning, multimodality, long context, and next generation agentic capabilities, arXiv preprint arXiv:2507.06261 (2025).
- [32] D. Guo, D. Yang, H. Zhang, J. Song, R. Zhang, R. Xu, Q. Zhu, S. Ma, P. Wang, X. Bi, et al., Deepseek-r1: Incentivizing reasoning capability in llms via reinforcement learning, arXiv preprint arXiv:2501.12948 (2025).
- [33] H. Touvron, L. Martin, K. Stone, P. Albert, A. Almahairi, Y. Babaei, N. Bashlykov, S. Batra, P. Bhargava, S. Bhosale, et al., Llama 2: Open foundation and fine-tuned chat models, arXiv preprint arXiv:2307.09288 (2023).