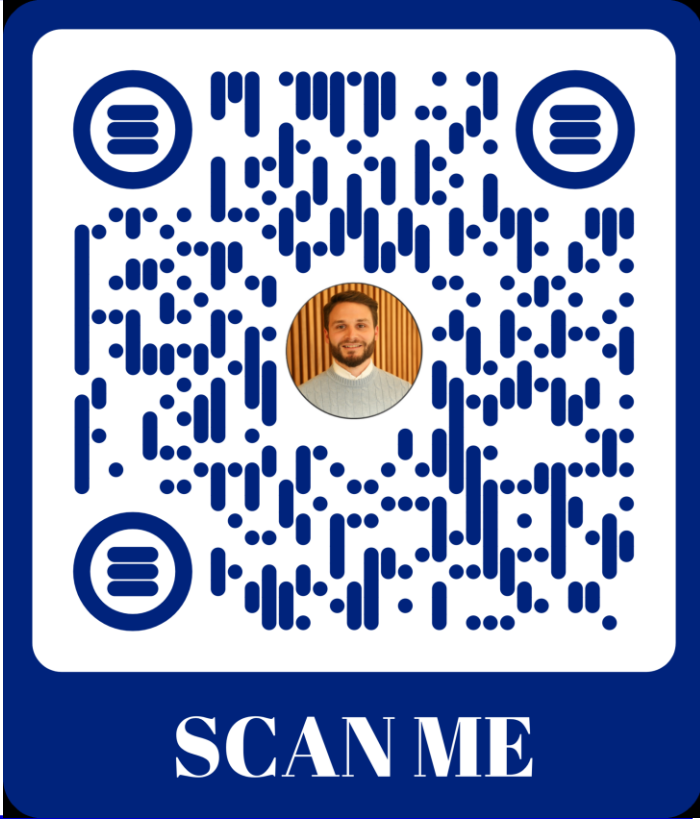


Leveraging Large Language Models for Causal Discovery: a Constraint-based, Argumentation-driven Approach

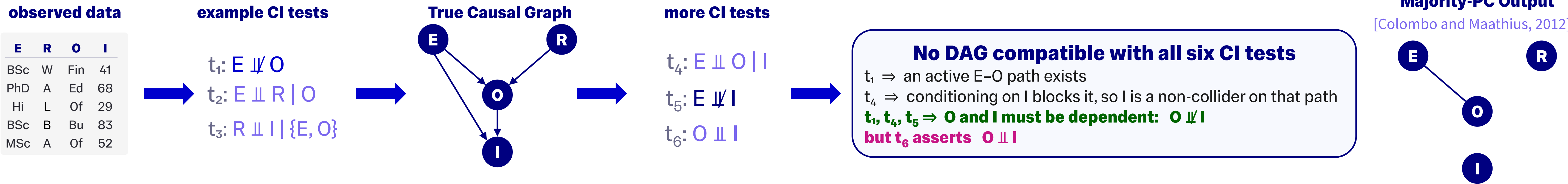
Zihao Li &
Fabrizio Russo



Motivation

Can we integrate semantic knowledge into data-driven causal discovery in a transparent, principled way?

Example



Background — Causal ABA and ABA-PC

Assumption-Based Argumentation framework $\langle L, R, A, \neg \rangle$

Assumptions $A \in L$

$A_{arr} = \{arr_{xy} | x, y \in V, x \neq y\}$ $A_{dag} = A_{arr} \cup \{noe_{xy}\}$
 $A_{ind} = \{(x \perp\!\!\!\perp y | Z) | Z \subseteq V, x, y \in V \setminus Z\}$ $A_{ds} = A_{dag} \cup A_{ind}$
 $A_{bp} = \{bp_{p|Z} | p \text{ an } x-y \text{ path}\}$ $A_{csl} = A_{ds} \cup A_{bp}$

Rules R

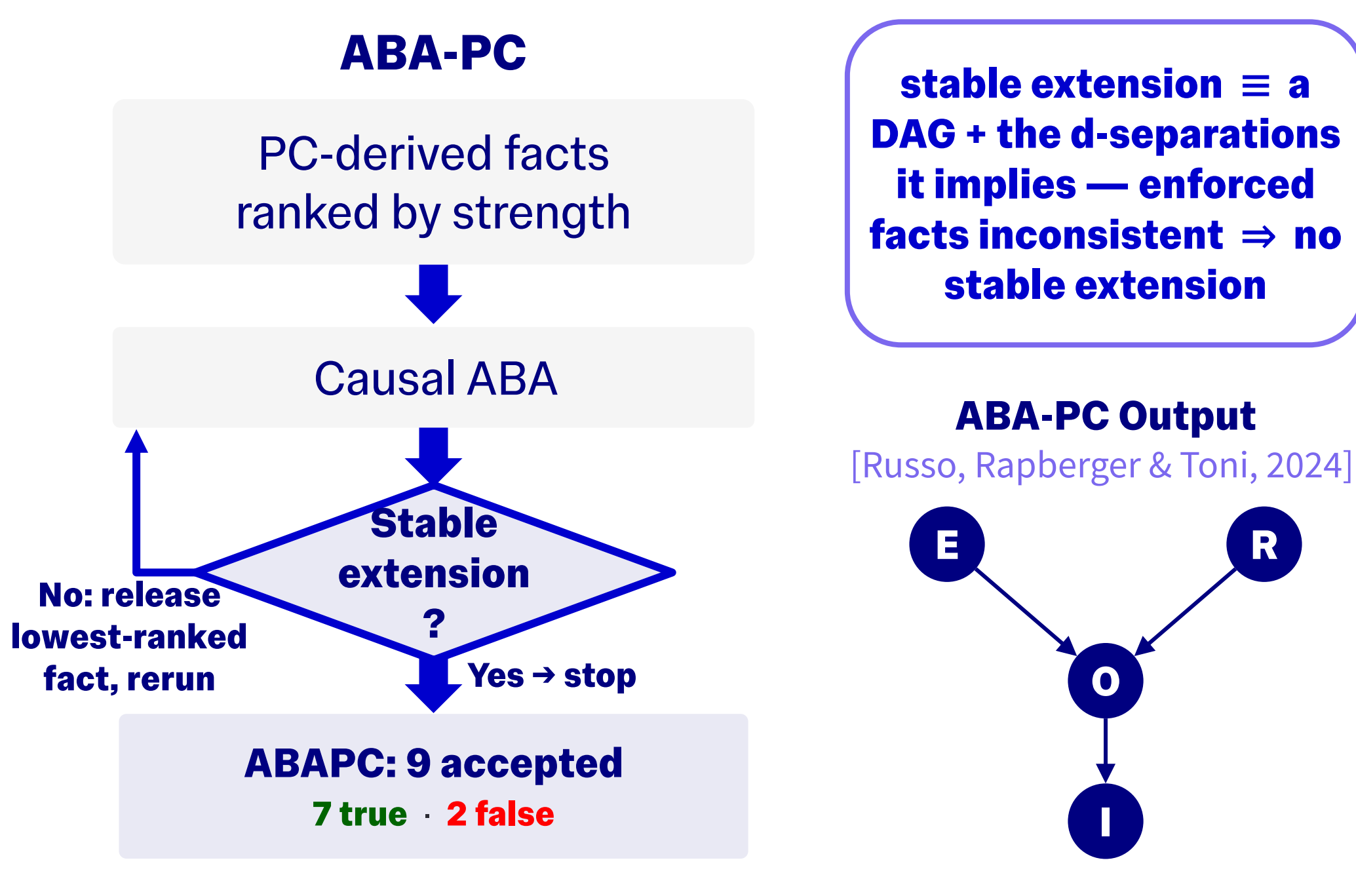
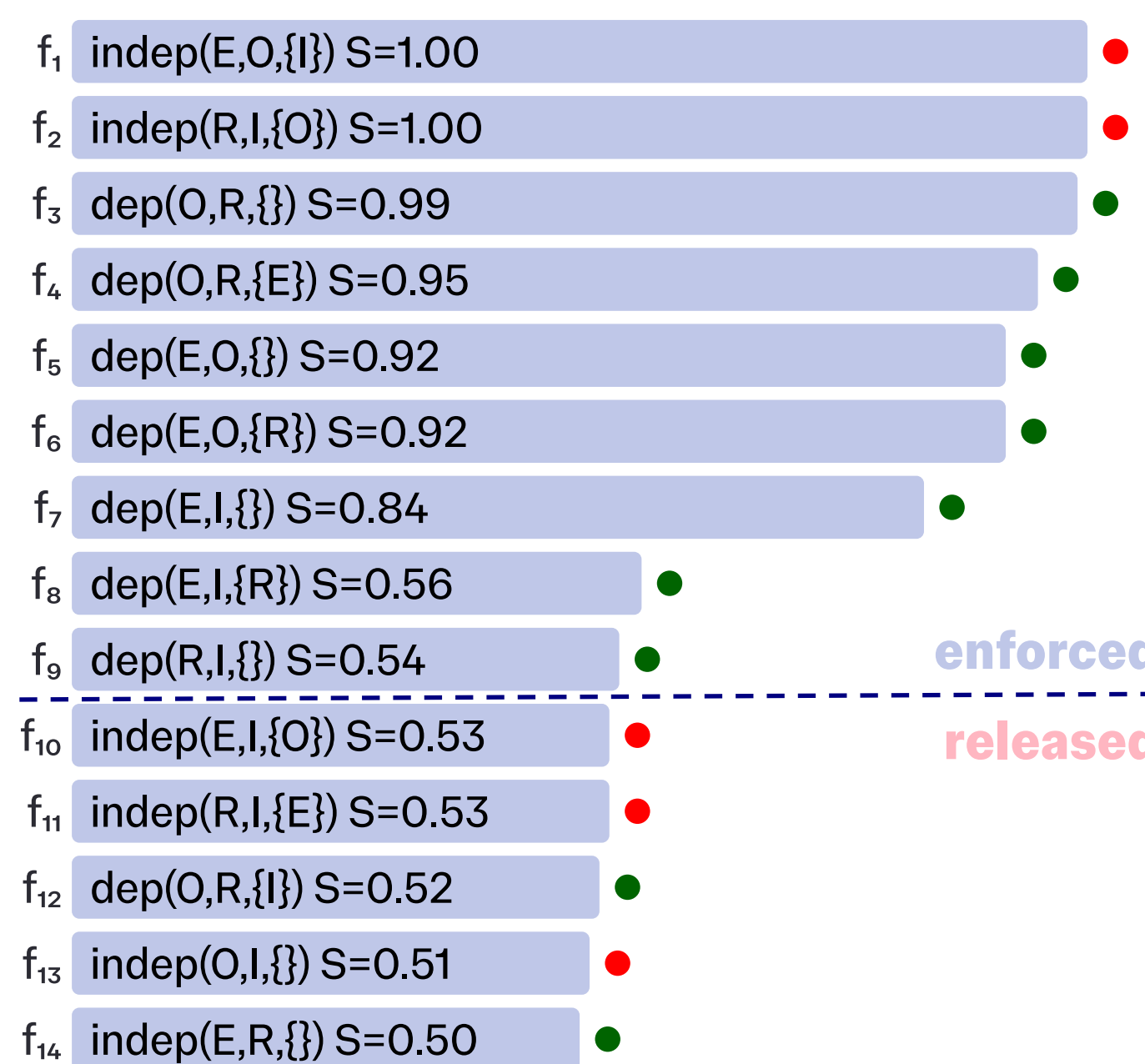
$R_{ds} = R_{dag} \cup R_{graph} \cup R_{act}$, $R_{ext} = R_{ds} \cup R_{xyz}$
 $R_{graph} : dpath_{xy} \leftarrow arr_{xy}$ $dpath_{xz} \leftarrow dpath_{xy}, arr_{yz}$
 $e_{xy} \leftarrow arr_{xy}$ $e_{yx} \leftarrow arr_{yx}$
 $R_{xyz} : (x \perp\!\!\!\perp y | Z) \leftarrow bp_{p_1|Z}, \dots, bp_{p_k|Z}$ $ap_{p|Z} \leftarrow p_t$
 $(x \perp\!\!\!\perp y | Z) \leftarrow \cdot (x \perp\!\!\!\perp y | Z) \leftarrow \cdot arr_{xy} \leftarrow$ **Facts: enter every extension**

Contraries $\neg : A \rightarrow L$

$\bar{a} \leftarrow b$, $a \neq b$, $a, b \in \{arr_{xy}, arr_{yx}, noe_{xy}\}$ one edge state per pair
 $\overline{arr_{x_1 x_{i+1}}} \leftarrow arr_{x_1 x_2}, \dots, arr_{x_{k-1} x_k}$ for $x_1 = x_k$ acyclicity
 $\overline{noe_{xy}} \leftarrow e_{xy}$ $(x \perp\!\!\!\perp y | Z) \leftarrow p_Z p_Z$ Z-active path $\overline{bp_{p|Z}} \leftarrow ap_{p|Z}$

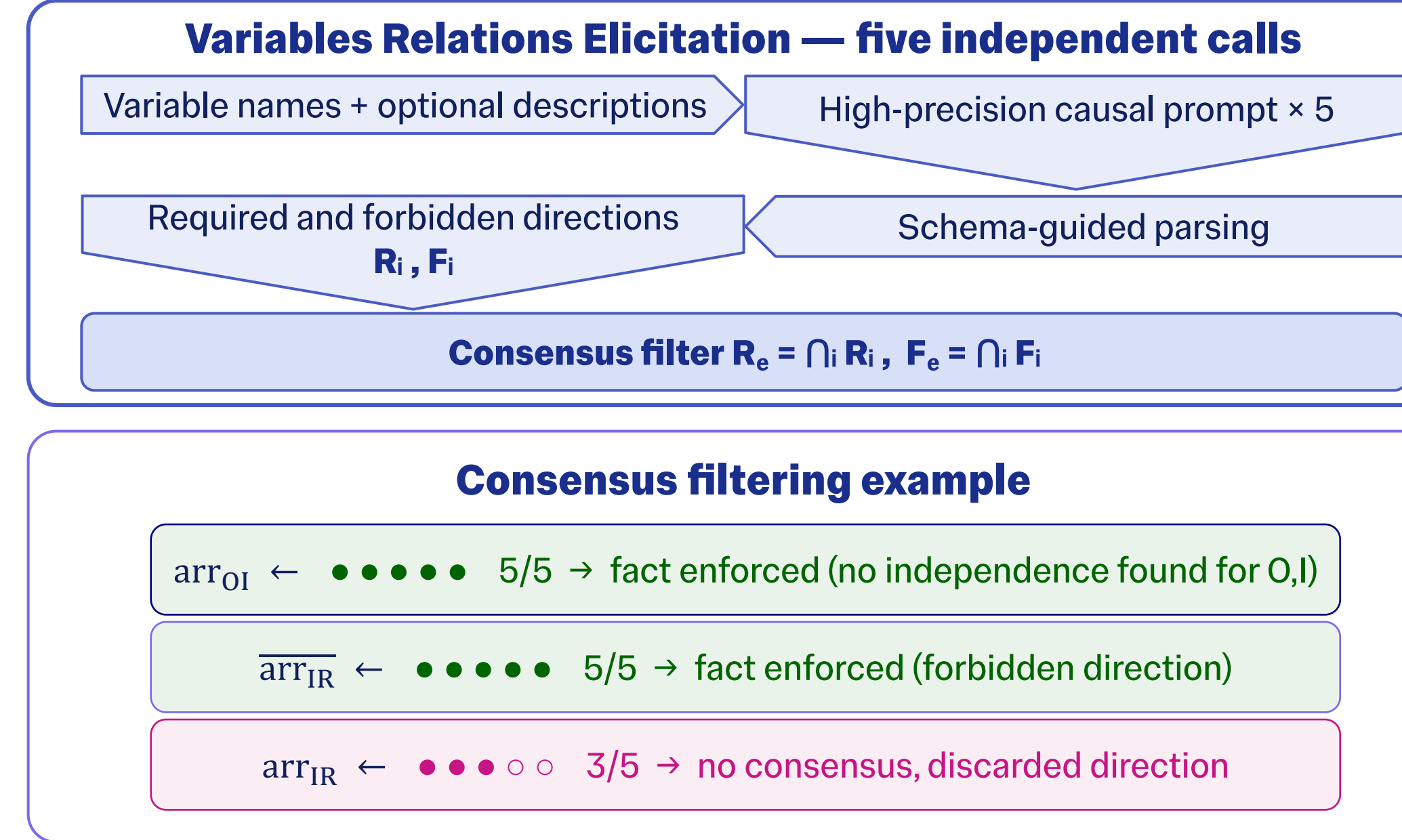
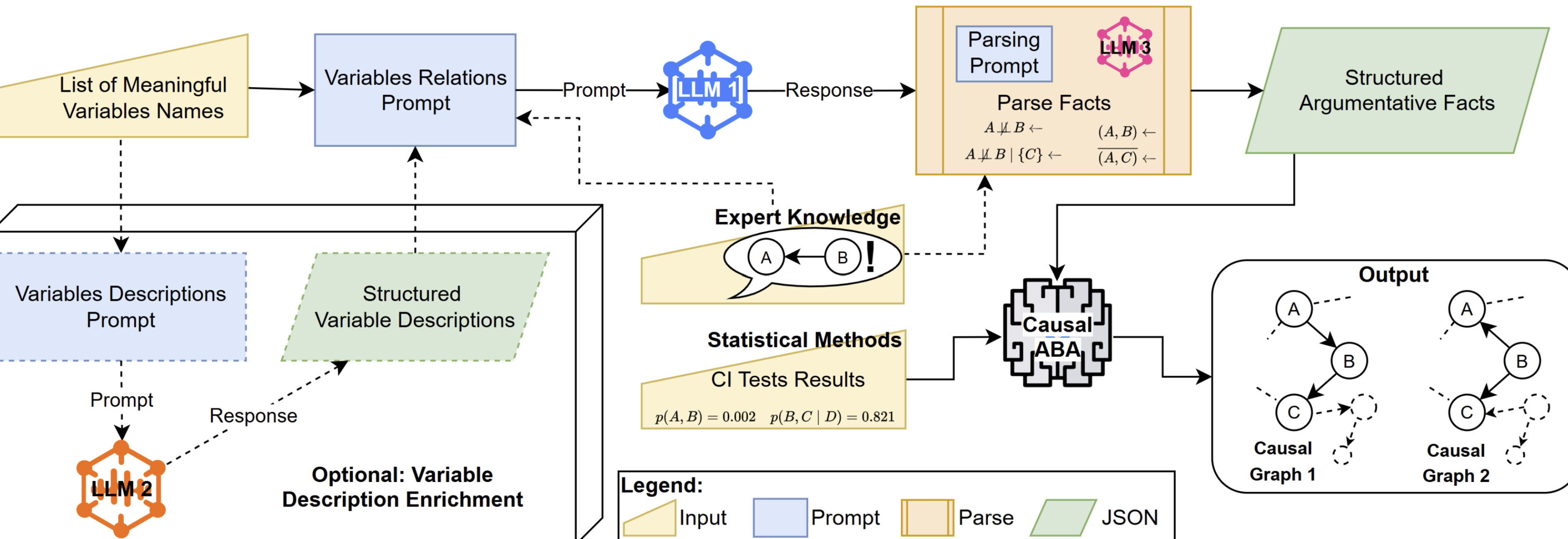
Stable Extension: set of assumptions that does not attack itself and attacks all assumptions outside the set.

14 CI tests, ranked by strength (p-value transformation)



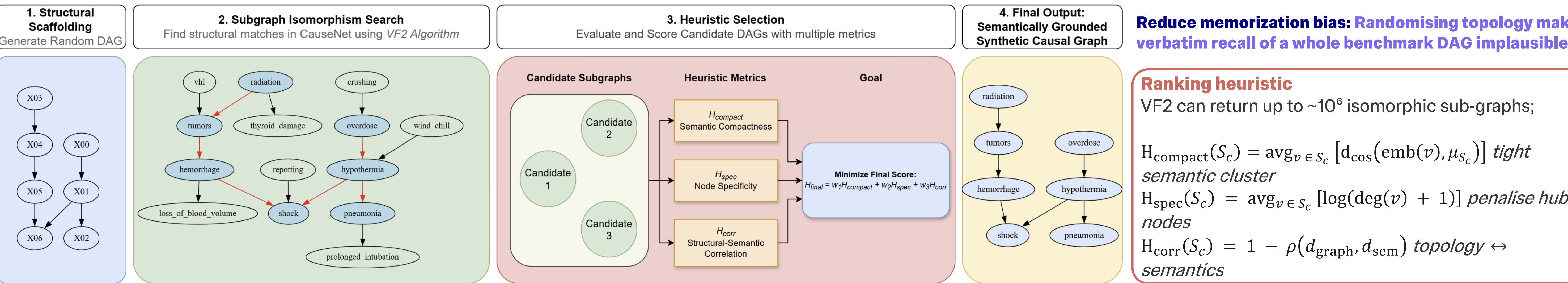
ABAPC-LLM: robust & transparent integration of LLM knowledge

Data constrains adjacencies; consensus semantic facts constrain directions. Integration through Causal ABA.



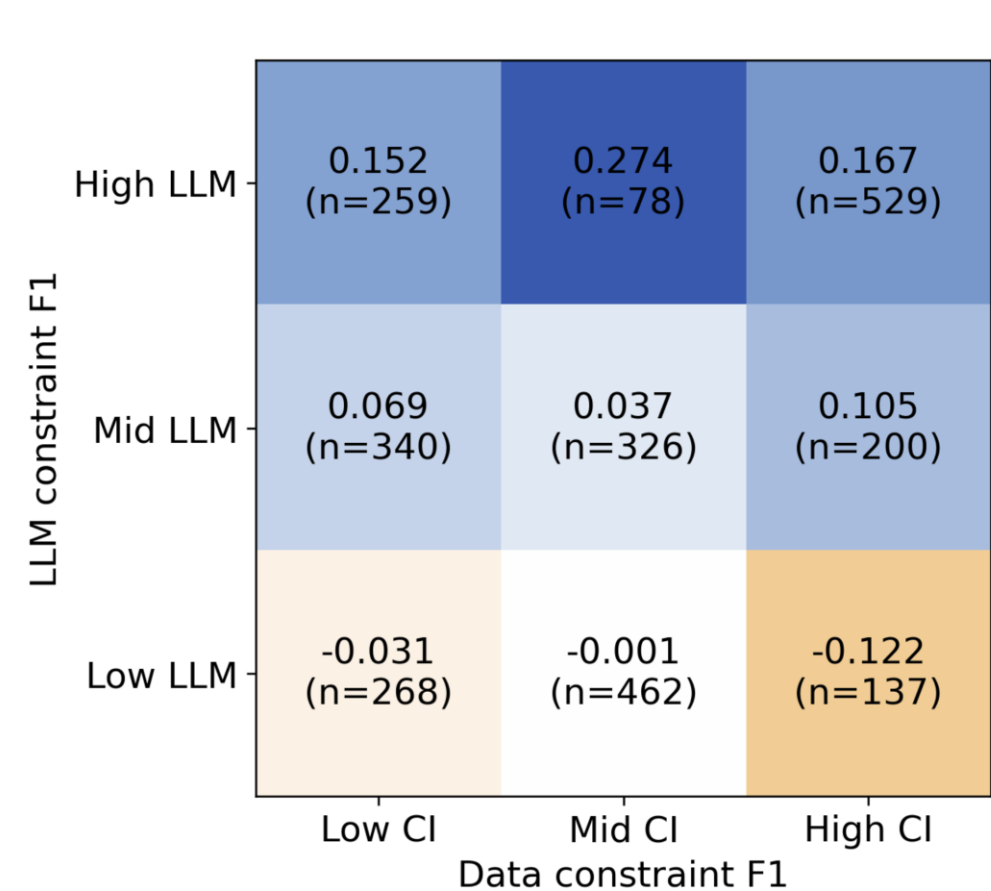
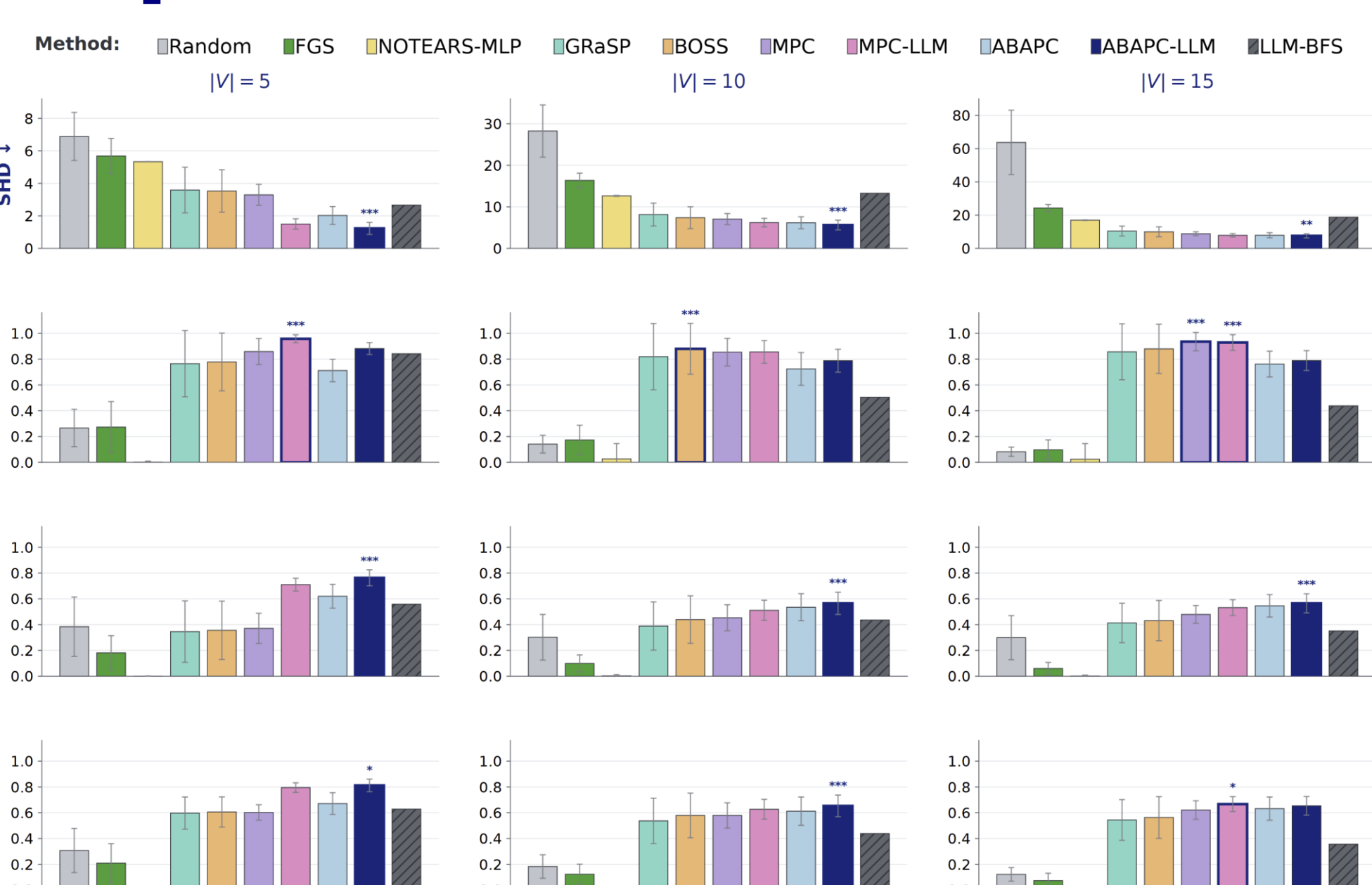
A required arrow $arr_{xy} \leftarrow$ is enforced only if x-y survives the reduced MPC skeleton; a forbidden arrow $\overline{arr_{xy}} \leftarrow$ blocks one orientation.
LLMs orient surviving adjacencies — they can never add an edge.

A novel benchmark for LLM-aided causal discovery

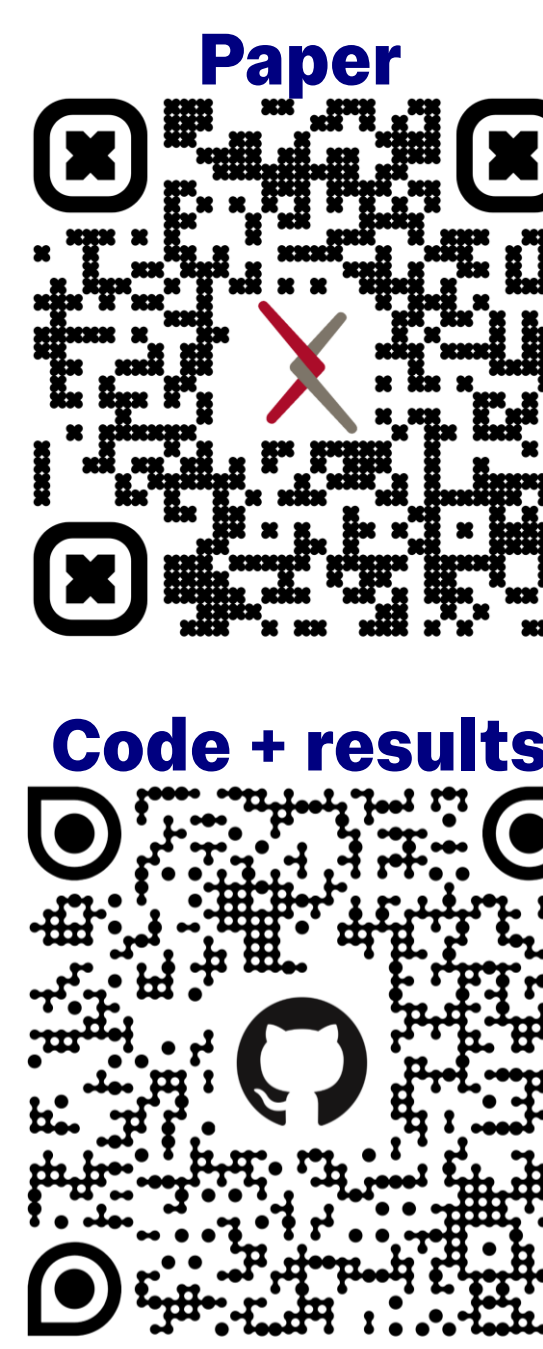


Empirical Results

54 DAGs · $|V| = 5/10/15$ · 5,000 samples · 50 reps · Wilks G^2 , $\alpha = 0.01$



- Key findings**
- Average SHD reduction**
- 57% vs LLM-BFS & -17.5% vs ABA-PC**
- Lowest SHD and highest recall at every graph size ($|V| = 5, 10, 15$); best F1 at 5 and 10.
 - SHD and recall gains are significant against the next-ranked method (** $p < 0.001$).
 - Semantic facts help most when both LLM and CI evidence are reliable.



IMPERIAL

uai2026