

Learning DiscoSCM under One Factual Event per Token: Stable Identity, Learnability Boundaries, and Semantic Tests

中文问题导读版

Anonymous Author(s)

Working paper, July 17, 2026

Guide version: LDSCM-GUIDE-ZH-2026-07-17-V4

Abstract

这份导读围绕一个总问题展开：**How to learn DiscoSCM?** 第一篇不尝试一次学习完整graph，而是在每个token 只有一个factual event 的条件下，研究fixed context 中的一个token-modulated outcome mechanism。V4 的核心不再是从 U 直接跳到Cauchy model，而是先分开target-population law P_U^{tar} 、actual realization $U = u_i$ 与learner belief Q_i ，再建立两条并行支路：合法factual evidence O_i 经Perception 与Abduction 形成一次 Q_i ；mechanism query x^Q 经query perception 后，与固定 Q_i 在Reduction 中汇合。本文只证明observed factual support 上的composed predictive-law consistency；E01、E05、E06 与P0/E07 则逐层区分composed law、fixed-selected-token surface、interface correctness 与architecture bottleneck。它们没有证明calibrated Q_i 、actual-token recovery、structural superiority 或Layer 3 causal identification。

1 怎样阅读这篇论文

一条不能倒转的逻辑链

target-population law $P_U^{\text{tar}} \Rightarrow$ actual $U = u_i \Rightarrow$ admissible evidence $O_i \Rightarrow$ focal belief $Q_i \Rightarrow$ fixed-belief query reduction $\Rightarrow$ learnability boundary $\Rightarrow$ tractable model and semantic tests。算法、likelihood 与robustness 都是后面的推论，不是论文的出发点。

英文主文把这条链写成definition、proposition、theorem 与evidence protocol。本导读不逐句翻译，而是用十二个问题帮助读者判断每一步究竟解决什么、又没有解决什么。

1. 总问题为什么必须是How to learn DiscoSCM?
2. 为什么第一步只学一个outcome mechanism?
3. 为什么Unit Selection Variable 不能删?
4. P_U^{tar} 、 $U = u_i$ 与 Q_i 为什么不是同一个对象?
5. 合法evidence O_i 与mechanism query x^Q 为什么必须分开?
6. Point 为什么会升级为distribution?
7. Gaussian 压力测试暴露了什么?
8. 为什么Cauchy 是natural tractable choice，而不是唯一真理?

9. conditional affine-in- U 算法怎样从P-A-R contract 中长出来?
10. one-shot data 真正能学到什么?
11. fixed abduction 何时才有causal 含义?
12. E01/E05/E06/E07 分别测了什么, 为什么不能pooling?

2 Q1–Q2: 总问题与最小slice

Q1: 这篇论文究竟在问什么?

不是“怎样做一个robust regressor”, 也不是“怎样写一个counterfactual algorithm”。总问题是: 如果DiscoSCM 是一套causal modeling theory, 它怎样从数据中被学习?

完整DiscoSCM 包含graph 与多个mechanisms, 也包含unit selection 和event noise。它还必须声明context、distribution consistency 与cross-world coupling。把这些都塞进第一篇, 会让“可学习”在不同段落里不断换含义。于是主文先固定一个不可再随意缩小的对象:

$$Y(x) = f_0(x, E; U). \quad (1)$$

f_0 是world-side mechanism; f_θ 才是learner family。即使well-specified 时存在 $f_0 = f_{\theta_0}$, one-shot data 也不因此唯一识别 θ_0 。

Q2: 为什么只学一个outcome mechanism 仍然不是普通regression?

因为它仍要回答outcome 属于哪个token、看不见的token code 怎样进入计算、factual evidence 怎样识别token、alternative query 怎样保持same-token, 以及factual data 到底监督了什么。

当前范围是fixed context、one factual row per token、一个predictors-to-outcome generator。它不声称已经学习多节点graph、一般distribution-consistency 或完整joint Layer 3 law。若这块地基还不清楚, “学习整个SCM”只会成为无法核验的口号。

V4 把factual row 精确写成

$$D_i = (O_i, X_i^F, Y_i^F) = (o_i, x_i, y_i). \quad (2)$$

O_i 是允许用来判断focal token 的evidence, X_i^F 是事实中真正发生的mechanism query。 Y_i^F 会跨records 监督共享learner parameters, 但当前slice 不把focal y_i 重新喂给abductor; 因此这里学的是 $P(U | O)$ 风格的response-blind belief, 不是 $P(U, E | O, X^F, Y^F)$ 风格的ex-post posterior。

还必须区分四个target:

1. factual composed law;
2. 固定 Q_i 后的alternative-query composed law;
3. simulator 固定actual u_i 后的token-specific response surface;
4. 另加assignment、support、stability、distribution consistency 与coupling 才能解释的causal target。

E01 测第二层, E05/E06/E07 才用truth vault 测第三层; 任何一层都不能自动升级成第四层。

3 Q3–Q4: 为什么需要(U), 又怎样计算它

Q3: 删掉Unit Selection Variable 会怎样?

若只写 $Y(x) = f_0(x, E)$, 稳定的between-token heterogeneity 只能被塞进event noise; “这个token的response mechanism”不再是模型对象; (x, x') 两个queries 也没有一个稳定selector 来保证仍属于同一token。

固定candidate u 后, 模型可以定义

$$Y_u(x) = f_0(x, E; u). \quad (3)$$

它不是每个token 各自拟合一条任意函数, 而是共享generator 的token-specific slice。这里的必要性不是“变量必须叫 U ”, 而是任何表达稳定heterogeneity、unit-indexed response 与same-token linkage 的理论都需要一个承担同等职责的primitive。

Q4: P_U^{tar} 、 $U = u_i$ 与 Q_i 为什么不是同一个对象?

P_U^{tar} 是declared target population 中token realizations 的composition; $U = u_i$ 是样本 i 对应的actual truth; Q_i 是learner 看见合法evidence 后, 对这个固定 u_i 形成的一整份epistemic belief:

$$Q_i(B) := q_\phi(U \in B \mid O = o_i) \approx P_{\text{tar}}(U \in B \mid O = o_i). \quad (4)$$

候选draws 是对同一个actual token 的hypotheses, 不是另一群真实units。若数据来自 $S^{\text{obs}} = 1$ 的selected sample, 直接学到的observed-population law 也不能默认等于target-population law; transport 需要独立assumptions。

计算上仍要区分semantic token 与representation。若actual token 是 $t^* \in \mathcal{U}_{\text{token}}$, 则

$$\iota : \mathcal{U}_{\text{token}} \rightarrow \mathbb{R}^d, \quad u^* = \iota(t^*). \quad (5)$$

本文把 ι 的单射性当作identity-encoding assumption。它不是one-shot data 推出的结论。为了简洁, 后文常压缩写成 $U = u^*$, 但这不表示物理token 就是向量。

尤其要区分两件事: token-to-code 可以是单射, code-to-response map 却不必要是单射。不同codes 完全可能诱导相同response law; 可逆重编码也可能不改变observable predictions。row index 只做bookkeeping, 不是generator input, 更不是一种新的selector ontology。

这里的“selection” 只有一个窄含义: U 选择token identity/code。它不是treatment/action A/T , 不是样本是否进入观察集的 S^{obs} , 也不是propensity。若写出 P_U , 它表示目标总体的token composition; Q_i 才是样本evidence 对focal token 的belief。

identity/code 与“响应等价类”也不是同一对象: 恢复相同response law, 不等于恢复physical token identity。

4 Q5–Q6: 为什么abduction 被迫出现

Q5: 世界和learner 的处境有什么不同?

world side 已经有actual code u^* , 所以它直接计算 $y = f_0(x, e; u^*)$ 。learner 看见 (O_i, X_i^F, Y_i^F) , 却没有 u_i label; 若仍要做token-modulated generation, 就必须从合法factual evidence O_i 形成focal-token belief。

因此不是把同一个raw tensor 随意复制成两个名字, 而是两条branch:

$$O_i \xrightarrow{P_F} z_i^F \xrightarrow{A_\phi} Q_i, \quad x^Q \xrightarrow{P_Q} z^Q, \quad (6)$$

最后在Reduction 中汇合:

$$\mathcal{L}_i(dy; x^Q) = \int p_\theta(dy | z^Q, u) Q_i(du). \quad (7)$$

这里的conditioning 只是learner-side inference, 不自动宣称world side 存在 $O \rightarrow U$ 或 $U \rightarrow O$ 。abduction 也不以显式population prior 或Bayes implementation 为定义。

做factual scoring 时, 只把 x^Q 设为实际发生的 x_i^F ; O_i 仍是另一份、职责不同的view。动作变量、sample inclusion 与hypothetical query 默认都不是unit evidence。

E06 说明, 只cache 一次还不够。若cached belief 在形成时已经读入与 U 无关的randomized action, evidence 与mechanism 的边界仍然会被污染。

Q6: 为什么point 会升级成distribution?

Point route 写作 $Q_i = \delta_{a_\phi(z_i^F)}$, 它断言剩余token uncertainty 可以忽略。one factual event 往往不足以支持这个强commitment, 于是learner-side result 可以升级为完整measure $Q_i(du) = q_\phi(u | z_i^F)du$ 。

$\tilde{U} \sim Q_i$ 是用于integration 的computational candidate, 不是actual identity 每次prediction 被重新抽样。若要比较同一个模拟token 在多种queries 下的outcome, 还必须只抽一次 $\tilde{u}_i^{(m)}$ 并复用; event noise 是否共享是另一项coupling decision。没有oracle token truth 时, point 与distributional routes 首先通过composed outcome law 接受比较, 而不是通过不可观测latent coordinates 自我认证。

5 Q7–Q8: Gaussian 压力测试与Cauchy 选择

Q7: Gaussian 错了吗?

没有。Gaussian 是最自然的distributional baseline。它承诺uncertainty 围绕finite-mean center 局部集中、远方candidates 被exponentially 压低、finite covariance 是合理summary, 并且Euclidean distance 对latent chart 有直接含义。真正的问题是: one factual observation 是否足以支撑这些认识论承诺?

主文提出六个desiderata: global openness、no forced moment commitment、清楚的location-scale semantics、affine stability、exact proper likelihood 与controlled residual influence。它们把“为什么不是默认Gaussian”变成一个可以被检查的设计问题。

Q8: 这些desiderata 为什么把我们带到Cauchy?

在coordinatewise symmetric stable location-scale candidates 中, Gaussian 是 $\alpha = 2$, Cauchy 是 $\alpha = 1$ 。Cauchy 具有polynomial tails、没有有限mean/variance, 并且affine sums 的scale 按absolute coefficients 精确相加。

我们写:

$$q_\phi(u | z_i^F) = \prod_j \text{Cauchy}\{u_j; m_{\phi,j}(z_i^F), \gamma_{\phi,j}(z_i^F)\}. \quad (8)$$

m 是location, 不是expectation; γ 是scale, 不是standard deviation。Cauchy 是在声明候选类与desiderata 后得到的natural tractable choice, 不是所有distributions 中被证明唯一的真理。latent chart 改变后“远近”也会改变。heavy-tailed candidate uncertainty 更不意味着任何token 都能产生任何outcome; outcome support 仍由generator 与event noise 决定。

6 Q9: 算法怎样从ontology 中长出来

图 1 把实现中的完整计算过程压到一张图里: 蓝色路径只从factual evidence 做一次unit abduction; 绿色路径允许同一token 接受多个新query; 金色模块把unit uncertainty 与event noise 解析合

成为outcome law。图下方同时区分训练时的factual support 与推理时的fixed-abduction protocol。

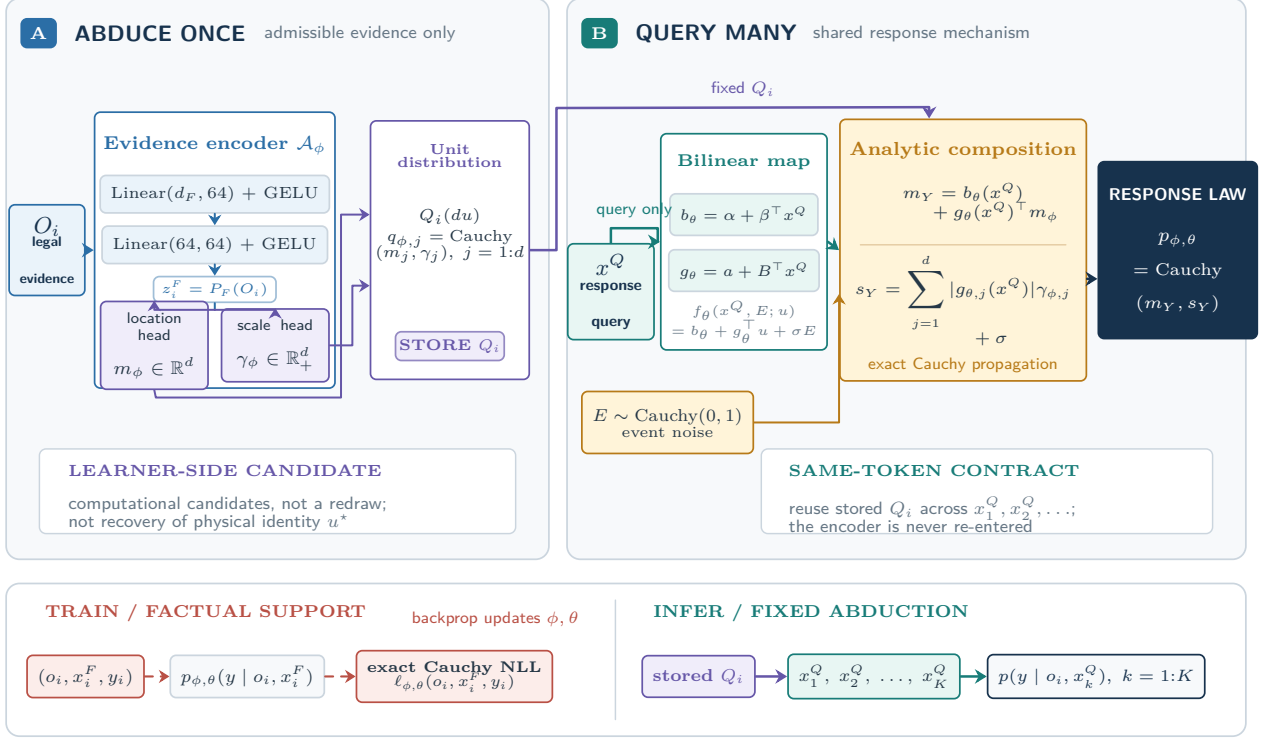


Figure 1: Learning DiscoSCM 的端到端神经网络与解析计算图。图中的64-64 GELU encoder、bilinear query head 和 $d = 4$ 是E01 benchmark 实例；location/scale 双头、conditional affine-in- U response、Cauchy 解析传播和“abduce once, query many”是更一般的模型结构。

令 $z_i^F = P_F(o_i)$, $z^Q = P_Q(x^Q)$ 。response 可以对perceived query 非线性，但对unit candidate 保持affine:

$$f_\theta(z^Q, E; u) = b_\theta(z^Q) + g_\theta(z^Q)^\top u + \sigma_\theta(z^Q)E. \quad (9)$$

raw bilinear model 是 P_Q 为identity、 $b = \alpha + \beta^\top x$ 、 $g = a + B^\top x$ 、 σ 为常数的special case。显式假设

$$E \sim \text{Cauchy}(0, 1), \quad E \perp\!\!\!\perp \tilde{U} \mid z_i^F, \quad (10)$$

candidate coordinates 在给定 z_i^F 后条件独立，且 $\gamma_j(z_i^F) > 0, \sigma_\theta(z^Q) > 0$ 。这时composed outcome law 仍为Cauchy:

$$m_Y(z_i^F, z^Q) = b_\theta(z^Q) + g_\theta(z^Q)^\top m_\phi(z_i^F), \quad (11)$$

$$s_Y(z_i^F, z^Q) = \sum_j |g_{\theta,j}(z^Q)| \gamma_{\phi,j}(z_i^F) + \sigma_\theta(z^Q). \quad (12)$$

因此learner 可以精确计算 $p_{\phi,\theta}(y \mid o_i, x^Q)$ ，无需latent Monte Carlo。软件contract 也应直接反映ontology:

```
abduce(factual_evidence) → UnitDistribution,
response_law_from_unit(unit_result, query) → ResponseLaw.
```

alternative query 只能复用第一个 Q_i 进入第二个operation, 不能重新调用evidence encoder。

这里的解析传播只证明tractability。它没有证明actual code、abduction distribution 或bilinear mechanism 已被one-shot data 唯一识别。Cauchy outcome NLL 的location score 对极端residual bounded 且趋近零, 是下游optimization robustness consequence, 而不是引入 U 的理由。

7 Q10: one-shot data 真正能学什么

训练只直接观察factual records (O, X^F, Y^F) 。令 $r_{\phi, \theta}(y | o, x) = p_{\phi, \theta}(y | o, x^Q = x)$ 。若true factual density $r_0(y | o, x)$ 位于模型类中且log risk 有限, 则proper log score 给出

$$\mathcal{Q}(r) - \mathcal{Q}(r_0) = \mathbb{E}[\text{KL}(r_0(\cdot | O, X^F) \| r(\cdot | O, X^F))] \geq 0. \quad (13)$$

所以能严谨声称的是: composed factual predictive law 在training support 上 P_{train} -almost surely Fisher-consistent。

不能从这条theorem 推出:

- actual u_i 或完整 Q_i 被恢复或校准;
- $(P_F, \mathcal{A}_\phi, P_Q, f_\theta)$ decomposition 唯一;
- token uncertainty 与event uncertainty 的scale split 唯一;
- 同一evidence o_i 下的alternative-query $p(y | o_i, x')$ 被factual data 识别;
- complete Layer 3 joint law 已被学习。

任意invertible latent recoding 都可以通过同步变换encoder 与generator 保持相同predictive law。在简单Cauchy sum 中, token scale 增加 δ 、event scale 减少 δ 也可保持total scale。更直接地, 两个models 只要在observed (o_i, x_i) support 上相同, 就拥有相同factual risk, 即使同一 o_i 下的alternative-query surface 完全不同。E05 观察到相近factual NLL 可以对应不同hidden token paths, 正是这条边界的empirical manifestation, 而不是theorem 的新证明。这些negative results 决定了“learnability”不能偷偷从predictive-law recovery 变成focal-belief、selected-token 或causal mechanism recovery。

8 Q11: fixed abduction 何时才是counterfactual

factual evidence $O_i = o_i$ 先产生一次 $Q_i = \mathcal{A}_\phi\{P_F(o_i)\}$ 。同一个token 面对 x' 的interface 是:

$$p_{\phi, \theta}(y | o_i, x') = \int p_\theta\{y | P_Q(x'), u\} Q_i(du). \quad (14)$$

不能用 x' 重新计算 Q_i , 因为那会用alternative query 重新识别一个potentially different token。若要比较同一个模拟token across actions, 还要从 Q_i 只抽一次 $\tilde{u}_i^{(m)}$ 并复用; 每个query 各自重抽即使marginal law 相同, 也已经换了模拟对象。

但fixed abduction 只关闭了re-recognition 这个软件/计算错误。要把它升级为causal Layer 3 statement, 还需要intervention semantics、token sufficiency、support/positivity、assignment 或confounding assumptions、mechanism stability、distribution-consistency 与event-noise coupling。 $Q_i(du)$ 或Cauchy 在latent-code space 的full support 不等于target population 中的treatment positivity / overlap; unit abduction 也不会自动修复sample-selection bias 或treatment confounding, 后两者需要独立的sampling/assignment 与identification argument。只问一个alternative

marginal law 时可积分声明的 E law; 若问多个 worlds 的 joint sample, 必须明确 $E_x, E_{x'}$ 是共享、独立还是采用其他 coupling。

Deep Structural Causal Models 在已声明 graph 中通常 abduct exogenous noise, 再执行 abduction-action-prediction [1]。当前 Learning DiscoSCM slice 首先 abduct 的是对 actual token code 的 epistemic uncertainty, 并把稳定 U 与 event E 分开。这不是 dominance claim; 完整模型可能同时需要 token abduction 与 exogenous-event abduction。

9 Q12: 什么实验才真正有辨别力

四个实验不是一个可以 pooling 的 leaderboard。E01 评分 evidence-indexed composed response law; E05/E06/E07 才在 simulator truth vault 中固定 actual u_i 后评分 selected-token response surface; 两者又都不是已识别的真实 causal counterfactual。先分清 target, 结果数字才有含义。

普通 regression benchmark 很适合检验 factual prediction, 却不能单独证明 token factorization。primary synthetic benchmark 的每个 token 只有一条 factual (w, x, y) record。simulator 会抽取 hidden actual code 来生成这一条 record; 但本次 E01 scoring oracle 并不把 actual code 当作 recovery target, 也不保存这个 code 的 21 个 realized potential outcomes。它只为每条 clean test evidence w 私藏 21 个 queries 上 composed conditional response law 的 location/scale。oracle truth 不得进入 fitting、validation、early stopping 或 hyperparameter selection。

当前落地的 bounded test 是 Y-shuffle hidden response-curve experiment。四个 worlds 是 matched、no heterogeneity、weak evidence 与 narrow factual query support。每个 world 有 6,000 个 token, 按 60/20/20 划分; evaluator 私藏 $[-1, 1]$ 上 21 个 query。只在合并后的 train+validation outcome pool 中, 用 SHA-pinned Sattolo subset derangement 破坏 $\rho \in \{0, .1, .2, .3, .4\}$ 的 factual links; outcome multiset 不变, selected links 全部移动, test 与 oracle 始终保持 clean 并被 fingerprint。

七个 models 是 Cauchy/Gaussian/deterministic Unit、matched direct Cauchy/Gaussian、varying-coefficient Cauchy 与 four-expert Cauchy mixture。primary capacity contrast 固定为 5,267 参数的 Cauchy Unit 对 5,216 参数、hidden dims (66,66) 的直接 Cauchy, 差距低于 1%。所有 inferential fits 固定使用 AdamW、learning rate 10^{-3} 、weight decay 10^{-5} 、batch 256、最多 300 epochs、patience 30。

Unit evaluator 对每个 clean test batch 只调用一次 `abduce(evidence)`, 保存 `UnitDistribution` 后复用于全部 21 个 queries; alternative query 不得重新进入 encoder。direct models 走 direct-query path。当前实现的 test-token abduction 只读 w , 不读取该 token 的 clean factual y ; 后者只用于 factual evaluation metrics。

formal primary endpoint 是 matched、 $\rho = .4$ 的 hidden curve-location error。对 confirmatory seed s , 定义

$$\Delta_s = \log \frac{\max(E_{\text{direct},s}, 10^{-12})}{\max(E_{\text{unit},s}, 10^{-12})}.$$

primary robustness gate 同时要求 $\text{mean } \Delta > 0$ 、two-sided 95% t interval 下界 > 0 , 以及 one-sided exact sign-test $p \leq .05$; ties 按冻结规则剔除。no-heterogeneity guardrail 把 Unit advantage 明确定义为 $1 - \text{gmean}(E_U/E_D)$, 其 95% upper bound 不得超过 10%。robustness confirmation 与是否允许归因于 token heterogeneity 分开报告, overall acceptance 要求两者都通过。

secondary outputs 包含 $\rho = .3$ 、relative-to-clean degradation、factual NLL/median error、pinball、PIT-ECE、coverage/width、hidden-grid CDF error, 以及 narrow-support 内外区 curve error。geometry、loss-family、capacity 三个 comparison families 各自在完整预声明 hypotheses 内做 Holm correction。non-degenerate Cauchy 没有有限 mean 与 variance, 因此不能把这些不存在的 moments 当作 metrics。

执行分三段：8-row mechanics smoke; seeds 11/22/33/44/55 的700-row development; 冻结source、dependency、config、analysis、gates 与output namespace 后，seeds 101–110 的1,400-row confirmatory。smoke/development 都不能promotion; confirmatory failure 或zero/reversal 也是正式结果，不能通过换一个official namespace 重开。

正式结果：primary 没有通过

canonical confirmatory matrix 已完成1,400/1,400 cells, zero failed fits, 冻结状态是not_confirmed。在matched、 $\rho = .4$ 的primary endpoint 上, $\bar{\Delta} = -0.00715$, 95% interval 为 $[-0.0496, 0.0353]$; Cauchy Unit 只在4/10 seeds 上优于direct Cauchy, one-sided sign-test $p = 0.8281$ 。所以mean > 0、interval 下界> 0、 $p \leq .05$ 三道门全部失败。这个实验没有确认Unit factorization 的robustness advantage。

no-heterogeneity guardrail 单独通过：apparent Unit advantage 的95% upper bound 是6.01%，低于10% ceiling。这只说明负对照中没有检测到material manufactured advantage; 因为primary 本身失败，它不能被解释成“heterogeneity 造成了优势”，也不能把overall verdict 翻成positive。

supportive secondary 中，matched、 $\rho = .3$ 的若干Cauchy Unit vs Gaussian Unit absolute prediction、curve 与calibration endpoints 在geometry-family Holm correction 后偏向Cauchy; 但clean-to-40%-shuffle relative degradation contrast 不支持Cauchy。更重要的是，这些是secondary results，而且matched DGP 本身按Cauchy 生成。因此它们不能确认Cauchy 是正确的epistemic geometry，也不能改变primary failure。

若direct Cauchy/Laplace 在tabular corruption 上表现好，这只支持heavy-tailed predictive objective 的某些factual robustness; 它不支持token-code recovery、unit factorization 或causal validity。结构claim 必须在matched hidden-query task 中获益，同时在no-heterogeneity、weak-evidence 与unsupported-query cells 中诚实报告无优势或失败。nonlinear misspecification、outcome contamination 与DSCM fixed-event-noise target 是后续独立extensions，不并入本次统计。

本次结果不是positive。即使结果为正，本实验也只可能支持synthetic fixed-abduction composed-law robustness; 它不证明恢复真实 U 、识别individual counterfactual、distribution-consistency 或完整Layer 3。

E05: 把target 改成fixed-selected-token surface 后，joint factorization 没过clean gate

E05 固定simulator-selected u_i 扫21 个queries。true mechanism frozen 时，learned abductor 相对finite-test evidence-conditioned reference 的最大normalized excess 是0.00729; 但abductor 与mechanism 一起从one-shot factual likelihood 学习时，matched clean normalized excess 为0.5307, no-heterogeneity error 为0.6808, 都超过0.20 ceiling。因此development_ready=false, 后续shuffle development/confirmatory 没有运行。相近factual NLL 不能排除错误的hidden-query basin。

E06: 只cache 一次还不够，形成cache 的evidence 也必须合法

在 $P(U | z, a) = P(U | z)$ 的randomized-action DGP 中，让abductor 只读token cue z , matched fixed-token curve MAE 从0.223607 降到0.030654, 并得到exact action invariance。但no-heterogeneity control 的改善更大，semantic-specificity guard 失败; frozen-mechanism 下learned Cauchy Q_i 的80% coverage 只有0.3822, 而exact reference 是0.7917。所以它只支持interface correction, 不支持 Q_i recovery、calibration 或corruption robustness。

E07/P0: Perception 与conditional response 值得继续，但尚未semantic-ready

108/108 formal fits 与48/48 nonranking diagnostics 表明raw bilinear model 确有representation 与function-class 限制; explicit query perception、对query 非线性而对 U affine 的response form 带来方向一致的改善。但invertible flow 没有胜过residual MLP, isolated paths、bad initialization basin 与no-unit spurious evidence-dependence 仍未过gate。正式状态是semantic_not_ready: 它只为下一版architecture 提供bounded motivation。

10 三个活跃论文身份与DSCM 对照

Paper identity	独立职责
Learning DiscoSCM	How to learn DiscoSCM; 当前聚焦一个token-modulated outcome mechanism 的end-to-end ML slice
Unit Mechanism Learning	shared unit-modulated generator、factorization、predictive-law consistency 与tractable family 的一般方法视角
Token Causation	row-as-token ontology、token response prediction 与singular/actual attribution 的理论边界

三篇可以互相引用与压力测试，但不能共享claim truth。一个theorem、experiment 或failure 进入另一篇时，必须重新说明assumptions、target 与evidence provenance。

与Deep SCM 的关键差别也应保持窄而准确：DSCM 的主要abducted object 是declared graph 中的exogenous noise；当前slice 的主要abducted object 是actual token representation 的candidate uncertainty。DSCM event exogenous variables 与这里的event E 不是低级版本，token abduction 也不是更一般的noise abduction；它们回答的是不同问题。

11 最终claim 边界

本文已经建立什么

一个区分 P_U^{tar} 、 $U = u_i$ 与 Q_i 的learning ontology；evidence perception、one-time abduction、query perception 与response reduction 的P-A-R contract；point/Gaussian/Cauchy 的统一learner-side semantics；conditional affine-in- U Cauchy law 的解析传播；observed factual support 上的log-score Fisher consistency；latent reparameterization、uncertainty split 与alternative-query underdetermination；以及逐层truth-isolated 的negative/diagnostic evidence。

本文没有建立什么

没有证明Cauchy 是唯一或universally best family；没有从 (O, X^F, Y^F) 唯一恢复actual token code、calibrated Q_i 、latent axes 或causal mechanism；没有由factual theorem 恢复alternative-query law；没有证明P-A-R architecture、invertible perception 或factorized model superiority；没有把fixed abduction 直接等同于Layer 3 identification。

读者判断这篇论文是否成立时，应依次追问： U 是否真是不可删除primitive？observation process 是否让 X 成为合理token evidence？Cauchy desiderata 是否清楚且可证伪？primary benchmark 是否真正隔离oracle truth？causal interpretation 是否逐项补足assumptions？任何一项不能靠更漂亮的NLL 或更大的neural network 跳过。

References

- [1] Nick Pawlowski, Daniel Coelho de Castro, and Ben Glocker. Deep structural causal models for tractable counterfactual inference. In *Advances in Neural Information Processing Systems*, volume 33, pages 857–869, 2020. URL https://proceedings.neurips.cc/paper_files/paper/2020/hash/0987b8b338d6c90bbedd8631bc499221-Abstract.html.

A 符号速查

符号	含义
t^*	actual discrete focal token
$u^* = \iota(t^*)$	world-side computational code
$\tilde{U}$	learner-side computational candidate, 不是identity redraw
$O_i = o_i$	focal token 的合法factual evidence
$z_i^F = P_F(o_i)$	abduction 使用的perceived evidence
Q_i	只形成一次并跨queries 复用的focal-token belief
x_i^F	factual outcome 实际发生时的mechanism query
x^Q	outcome mechanism query
$z^Q = P_Q(x^Q)$	perceived query representation
E	token 与query 固定后仍变化的event randomness
f_0/f_θ	world mechanism / learner family
$p_{\phi,\theta}$	candidate-unit uncertainty 与event noise 合成后的predictive law

若只记住一个检查规则：**训练数据只监督observed (o_i, x_i) factual support**；任何关于同一 o_i 下**alternative query**、**actual code**、**calibrated Q_i** 或**cross-world coupling** 的更强结论，都必须指出额外结构与证据来自哪里。