



智能体时代的小样本学习

结构化上下文：连接模型泛化与智能体适应的新范式

姚权铭

清华大学 电子工程系

2026.08

内容概要

① 时代背景

AI 任务从静态预测走向智能体执行，小样本学习面临新的挑战

② 核心思路

有限监督不等于有限信息：结构化上下文连接模型泛化与智能体适应

③ 研究探索

快速泛化机制、高效适配方法、可靠外推方法与药物安全应用

④ 未来方向

从静态样本关系走向动态交互经验

主线：模型泛化 → 结构化上下文 → 智能体适应

智能体时代：AI 从预测答案走向完成任务

AI 从一次性输入输出，走向多步骤、交互式、长期执行的复杂智能体任务

任务形态演进

分类预测任务



应用示例：图像分类

任务特点：识别判断，输出确定类别

内容生成任务



应用示例：代码生成

任务特点：基于指令，生成结构化内容

多轮交互任务



应用示例：终端助手

任务特点：多轮对话，理解上下文意图

工具执行任务



应用示例：实验操作

任务特点：调用工具，完成具体操作

系统协同任务



应用示例：多机协同

任务特点：协同多系统，实现整体最优

长期代理任务



应用示例：大模型控制卫星

任务特点：长期规划，自主决策与持续执行



2012



2021



2022



2023



2024



2025

任务形态变化：从“看见一个答案”走向“完成一个过程”

为什么智能体更需要小样本学习？

任务从静态预测走向复杂执行，可获得的高质量显式监督快速下降。

简单任务：图像分类



猫

- 监督：类别标签
- 特点：简单，易获取

亿级数据

复杂任务：智能体优化化学反应条件



特点：多步决策、反馈延迟、调试成本高

图像分类

JFT-300M

3.75亿标签

代码任务

Project CodeNet

1400万样本

终端助手

Android in the Wild

3万指令

实验操作

RLBench

约1万演示

多机协同

POGEMA

3376实例

自主导航

Habitat ObjectNav

216场景

高质量显式监督快速减少

亿级

万级

千级

百级

如何让智能体在有限、稀疏、延迟的监督下，从少量经验中可靠学习与适应？

智能体时代的小样本学习：有限监督不等于有限信息

样本少了，不代表任务信息少了

标签只是冰山一角



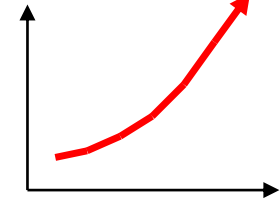
可见的标签很少，真正丰富的信息藏在任务上下文中。

例子：智能体优化化学反应条件



实验记录

产率



标签之外的信息很多

温度、浓度、pH、试剂关系、
实验过程、安全约束

显式监督很少

最终产率、成功/失败结果、
少量历史实验记录

状态 当前温度、压力、浓度、pH、液位、搅拌转速、仪器读数、环境参数

关系 试剂、溶剂、催化剂、配比、加料顺序与产率、选择性和副产物之间的依赖关系

过程 加样、升温、恒温、搅拌、取样、观察、调参、记录、终止反应与后处理

智能体时代关注少量标注样本和交互经验，分散信息如何被有效组织和利用。

研究思路：以结构化上下文为突破

人类认知启发

观察：人不是逐条记忆

启发：经验需要组织成图式

外部分散信息



复杂经验不是被完整记住
而是被抽象成可复用结构

研究思路：从任务数据中归纳结构化上下文

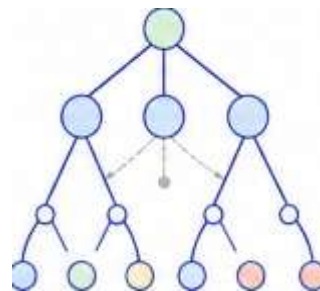
任务信息



状态 关系 过程
约束 反馈 历史案例

标签之外有很多信息，
但只是分散存在。

结构化组织



对象 状态 操作 结果
依赖 顺序 因果 规则

抽取结构单元，
建立关系与约束。

形成化上下文



任务图 / 状态机 / 关系子图
证据链 / 交互轨迹

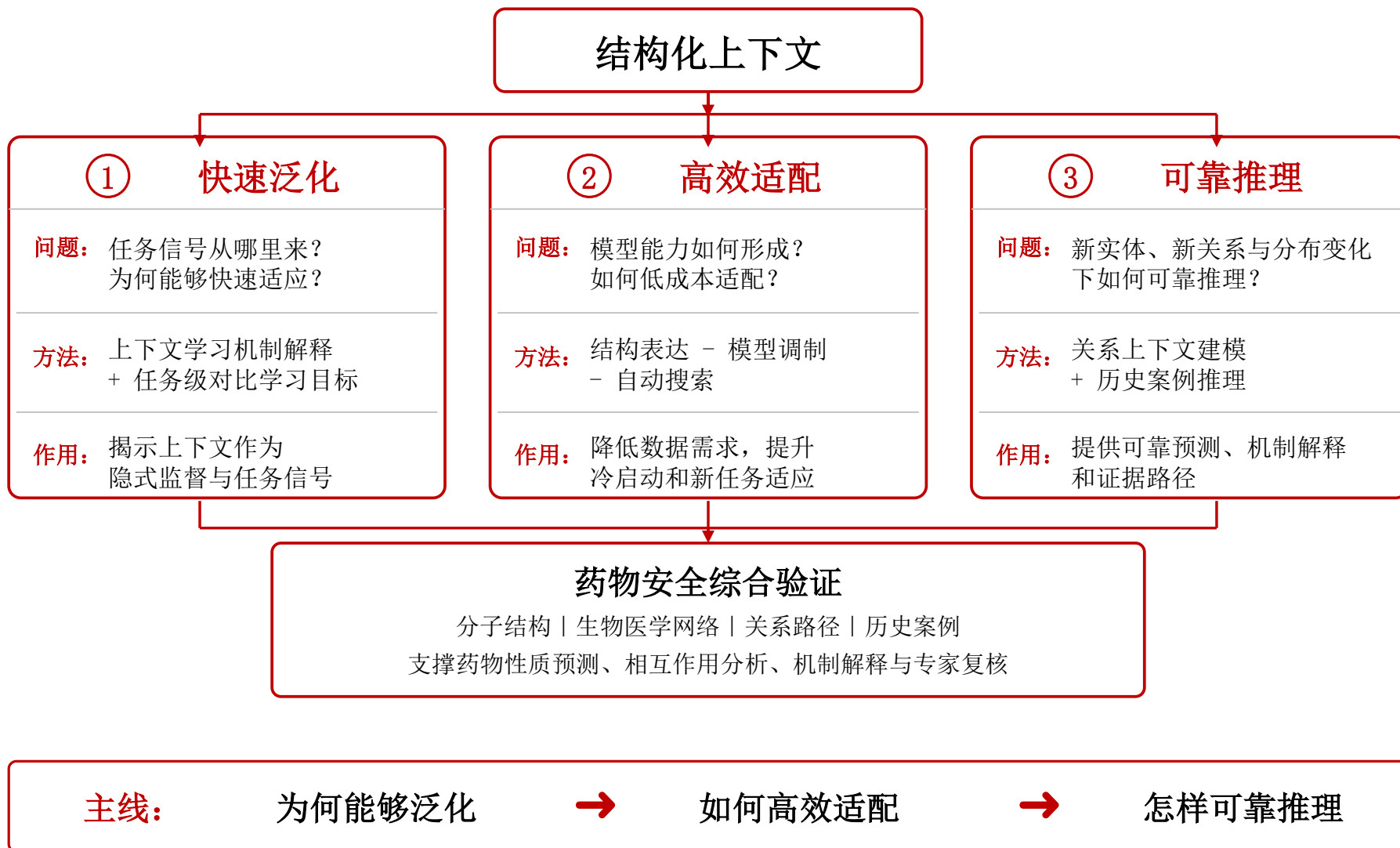
可学习 可复用 可推理

形成模型可用的结构，
支持学习、复用与推理。

分散信息 → 结构化组织 → 模型可用能力

结构化上下文使分散经验变得可学习、可复用、可推理。

结构化上下文：连接模型泛化与智能体适应



结构化上下文作为隐式监督 — 问题与思路

问题：有效果但机制不清晰

GPT-3报告开创性地发现大语言模型具有小样本**上下文学习能力**，引爆了通用大语言模型的时代

Demonstrations

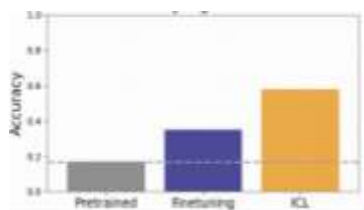
Circulation revenue has increased by 5% in Finland. \n Positive
Panostaja did not disclose the purchase price. \n Neutral
Paying off the national debt will be extremely painful. \n Negative
The acquisition will have an immediate positive impact. \n _____

Test input

LM

(无需更新参数)

Prediction Positive



预训练模型直接预测
vs
利用小样本数据训练微调
vs
利用小样本数据上下文学习

为了分析上下文学习的快速适应能力，关键问题是：

- 任务信号从哪里来？
- 上下文如何替代部分显式监督？
- 这种适应能力是否稳定可迁移？

Language models are few-shot learners

T. Brown, B. Mann, N. Ryder, M. Subbiah, JD. Kaplan, P. Dhariwal, A. Neelakantan, et al.

被引用次数: 71781

Advances in neural information processing systems, 2020 · proceedings.neurips.cc

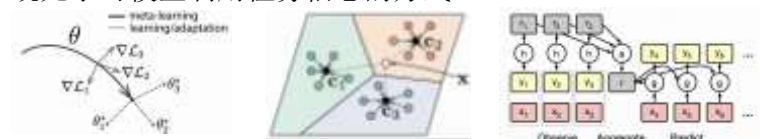
思路：从元学习视角看上下文结构

人类具有从少量样本中泛化的快速适应能力



少量样本	共享局部结构	任务概念
江、河、湖、海	共享“氵” 局部结构	水相关语义
打、拍、推、拉	共享“扌” 局部结构	手部动作
妈、吗、码、骂	共享“马” 局部结构	读音线索

元学习即使神经网络学习如何学习：不只是记住样本，而是可以识别、生成、解析结构并迁移到相关概念。在小样本任务中获得接近人类的系统泛化能力（Nature 2023）。几类传统元学习模型利用任务信息的方式：



上下文学习展现了新一种方式：prompt输入中存在结构性任务知识，利用模型的前向传播完成任务级监督信息的学习与适应、预测：

$[x_1, f(x_1); x_2, f(x_2); x_3, f(x_3); x_4]$

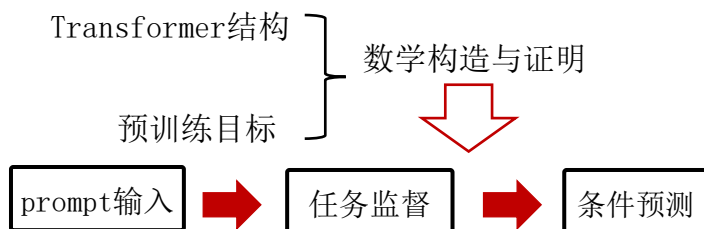
Transformer $\rightarrow f(x_4)$

从少量样本中理解上下文为何能够支持快速泛化

结构化上下文作为隐式监督 — 创新与成果

创新点：把上下文中的任务结构转化为元学习监督信号

机制认识：Transformer通过预训练学会利用结构化上下文学习并预测



- 将上下文学习解释为隐式元学习过程，任务内部样本关系提供适应信号——**统一上下文学习与元学习理论框架**

$$\hat{\mathbf{y}}^{(q)} = g(\mathbf{x}^{(q)}, \mathcal{D}) \quad g_M(\mathbf{x}^{(q)}, \mathcal{D}; \theta_M) = \text{TF}_M(Z_0; \theta_M)$$

- 证明上下文对学习算法的表达能力超越已知种类的元学习模型——**上下文适应机制的理论可能性**

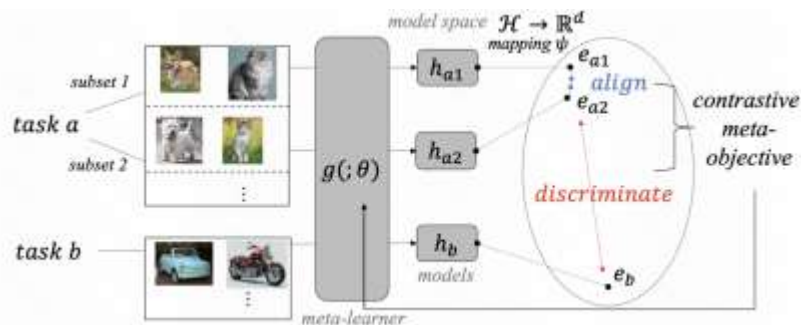
Theorem 2.1. $\forall g \in \{g_{gd}, g_{sim}, g_{prt}, g_{aw}\}, \forall \theta \in \mathbb{R}^{|\theta|}, \exists M \in \mathbb{N}^+ < \infty, \exists \theta_M \in \mathbb{R}^{|\theta_M|}, g_M(\mathbf{x}^{(q)}, \mathcal{D}; \theta_M) = g(\mathbf{x}^{(q)}, \mathcal{D}; \theta).$

- 证明任务结构信息带来泛化能力保障——**上下文学习中的泛化性优势**

Theorem 1. $\forall p(\tau), U_{p(\tau)}(\theta_{\mathcal{L}_u}^*) = \min_{\theta} U_{p(\tau)}(\theta), \text{ where } \theta_{\mathcal{L}_u}^* = \arg \min_{\theta} \mathcal{L}_c(g(\cdot; \theta), p(\tau)).$

机制利用：将多种方法迁移到元学习层次改进上下文学习的预训练过程

- 元对比学习提高泛化能力



针对ICL模型：通过prompt注入探针token得到任务表征 $\mathbf{e} = g([x_1, f(x_1); x_2, f(x_2); u])$

对齐：拉近同一个任务不同子集的任务表征间的距离

+

区分：拉远不同任务的任务表征间的距离

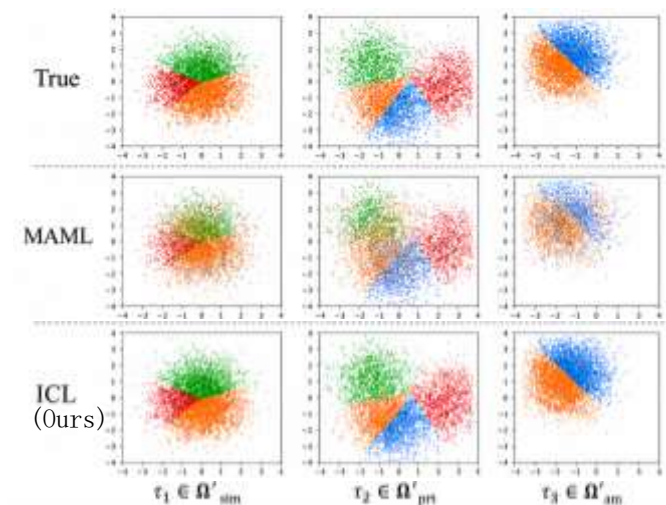
元对比目标（+原目标=改进预训练）

- 元课程学习加速模型收敛
- 元元学习提高快速适应能力
- ...

结构化上下文作为隐式监督 — 效果与影响

理论验证

上下文学习 **算法表达能力** 覆盖并超越传统元学习模型



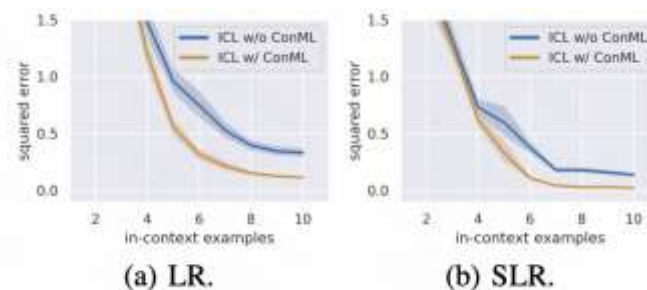
- 单个上下文学习（ICL）模型能够同时适应三类需要不同的学习算法（sim/prt/am）的任务，超越传统元学习（MAML）模型

方法效果

通过改进上下文学习模型的训练目标，增强快速适应与泛化能力

结论：小样本学习的关键不仅是样本数量，还包括上下文中的任务结构

- 在多种形式任务中，提升给定样本数量下的预测性能，适应更加快速



- 小样本图像分类性能在SOTA上下文学习模型上进一步提升，扩展泛化极限

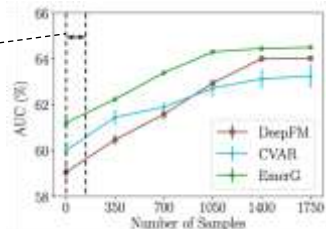
Objective	miniImageNet		tieredImageNet	
	5-way 1-shot	5-way 5-shot	5-way 1-shot	5-way 5-shot
-	96.15 $\pm$ 0.10	98.57 $\pm$ 0.08	95.41 $\pm$ 0.10	98.06 $\pm$ 0.10
w/ ConML	97.03 $\pm$ 0.10	98.92 $\pm$ 0.08	96.56 $\pm$ 0.09	98.23 $\pm$ 0.05

结构化上下文作为适配信号 — 问题与思路

问题：新任务下模型难以低成本适配

新任务/冷启动场景中**反馈少**；不足以支撑模型
重新训练

新用户交互数量
新加入电商平台的用户与各类商品的历史交互数量<60条；
推荐系统无法准确捕捉这些用户的商品偏好



重新标注数据训练，适应**成本很高**

已有样本稀缺

新任务样本标注成本极高



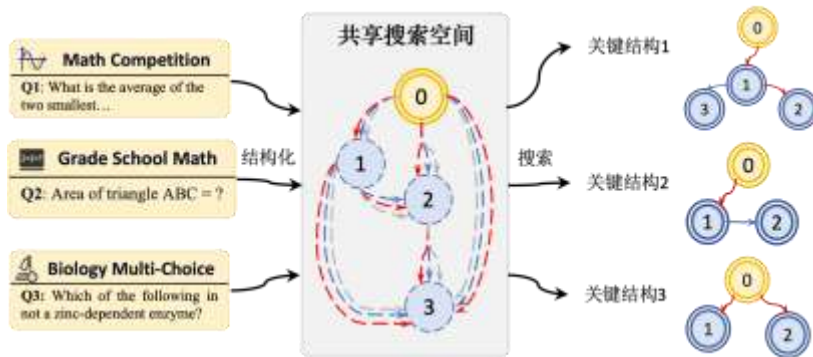
关键问题：模型需要快速改变，但显式监督
不足以支撑模型重新学习

思路：从模型重训转向结构化适配

人们一般不会为特殊路况重新造车，而是适配关键部件



基础模型共享 + 关键结构搜索

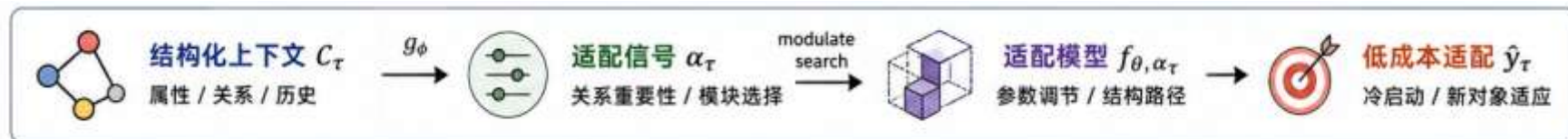


多个任务共享基础模型，无需重新训练
结构化上下文刻画任务差异，搜索轻量结构

从少量数据中识别任务结构，实现低成本适应

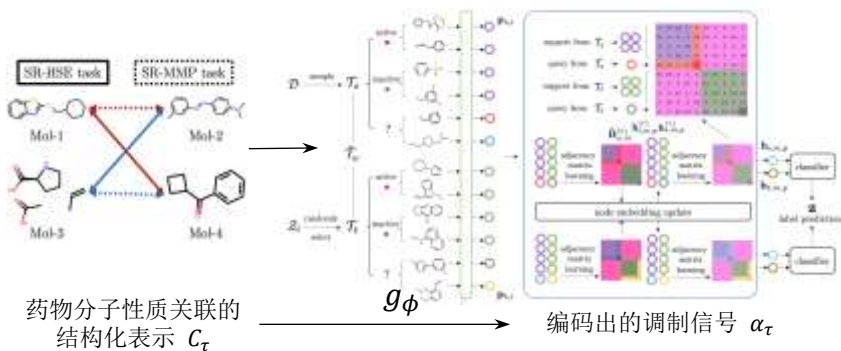
结构化上下文作为适配信号 — 创新与成果

创新点：将上下文结构转化为模型调制与搜索信号



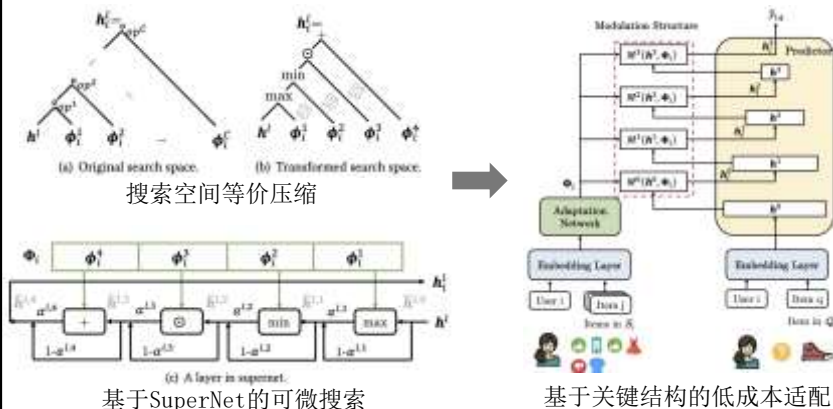
$$C_T \xrightarrow{g_\phi} \alpha_T, \quad \hat{y}_T = f_{\theta, \alpha_T}(x) \quad \text{结构化上下文经过适配函数 } g_\phi \text{ 生成任务适配信号 } \alpha_T, \text{ 用于调节基础模型 } f_\theta, \text{ 实现低成本预测。}$$

结构编码为适配信号：利用结构化先验区分任务相关关系，并基于此将结构化上下文编码为模型适配信号



- 通过自适应关系图学习显式区分不同任务下的分子关联拓扑，实现性质驱动的结构化关系
- 通过元学习解耦跨任务共享先验与任务特异参数更新，将分子结构化关联编码为对分子表征的适配信号

模型低成本适配：基于适配信号搜索关键结构和参数，实现模型在新任务上的低成本适配



- 基于算子群内一致性与群间置换不变性，将指数级模型搜索空间理论等价压缩为固定规模搜索空间
- 在等价转化后的紧致SuperNet中，以连续化结构权重实现端到端可微关键结构搜索

提出结构化上下文引导的模型高效调制算法，形成面向新任务的低成本适配能力

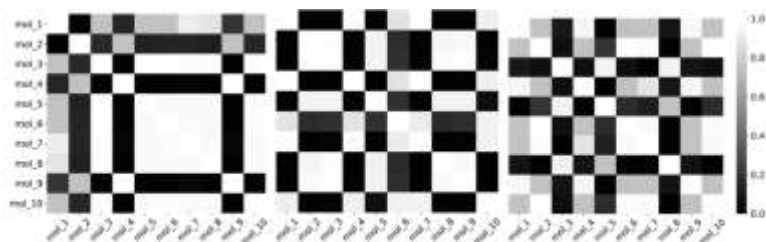
结构化上下文作为适配信号 — 效果与影响

效果：面对新任务，仅调整任务相关结构，降低数据需求与适配成本

在新任务小样本下ColdNAS效果显著

Dataset	Metric	DropoutNet	MeLU	MetaCS	MetaHIN	MAMO	TaNP	ColdNAS-Fixed	ColdNAS
MovieLens	MSE	100.90 _(0.70)	95.02 _(0.03)	95.05 _(0.04)	91.89 _(0.06)	90.20 _(0.22)	89.11 _(0.18)	91.05 _(0.13)	87.96_(0.12)
	MAE	85.71 _(0.48)	77.38 _(0.25)	77.42 _(0.26)	75.79 _(0.27)	75.34 _(0.26)	74.78 _(0.14)	75.65 _(0.30)	74.29_(0.20)
	nDCG ₃	69.21 _(0.76)	74.43 _(0.59)	74.46 _(0.78)	74.69 _(0.32)	74.95 _(0.13)	75.60 _(0.07)	75.11 _(0.09)	76.16_(0.03)
	nDCG ₅	68.43 _(0.48)	73.52 _(0.41)	73.45 _(0.56)	73.63 _(0.22)	73.84 _(0.16)	74.29 _(0.12)	73.89 _(0.12)	74.74_(0.09)

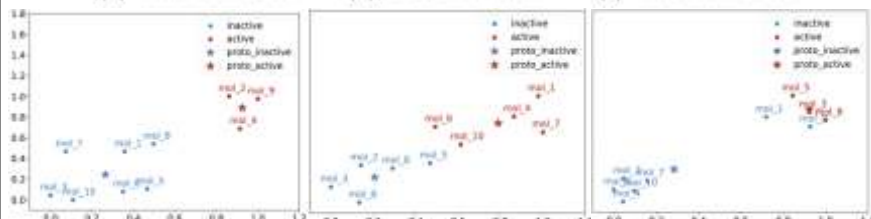
PAR实现了对分子间关联精准的结构化建模，
对小样本分子性质展现出卓越的预测效果



(a) The 10th task.

(b) The 11th task.

(c) The 12th task.

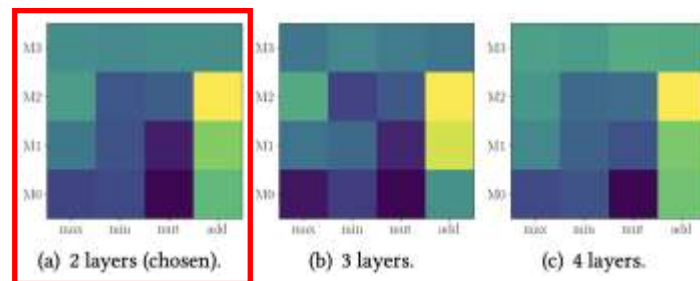


(a) The 10th task.

(b) The 11th task.

(c) The 12th task.

ColdNAS实现了对关键轻量结构的识别和
低成本适配

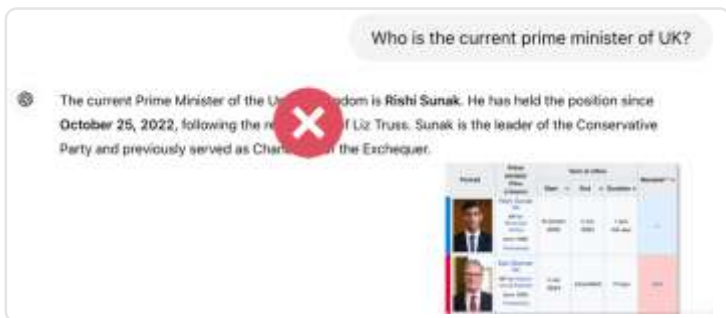


Clock time (min)		MovieLens	BookCrossing	Last.fm
ColdNAS	TaNP	15.5	44.2	4.1
	Search	16.2	45.5	3.9
ColdNAS	Retrain	12.7	35.5	3.5

结构化上下文作为推理证据 — 问题与思路

问题：新实体新关系下预测解释困难

在真实世界中，新实体、新关系和事件状态不断出现。
仅给出一个预测结果，难以判断：为什么可信？证据在哪？
失败原因是什么？



大模型幻觉在外推任务中有很严重的后果

- 过时/缺失事实会被包装成“确定答案”
- 误导事实引用、风险判断与后续决策，并形成伪证据链

如何搜集并组织可靠的推理证据，避免“伪证据”的误导，是解决外推任务的关键

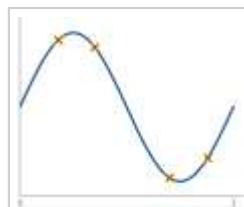
Detecting hallucinations in large language models using semantic entropy

Sebastian Farquhar, Jannik Kossen, Lorenz Kuhn & Yarin Gal

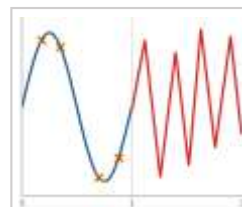
Nature 630, 625–630 (2024) | Cite this article

思路：利用结构化上下文组织证据

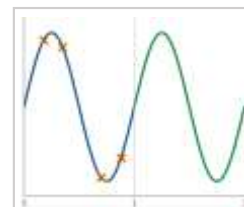
利用上下文进行函数外推



已知观测：[0, 1]



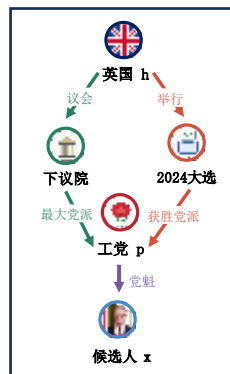
无上下文：杂乱外推



有上下文：三角函数约束

利用结构化上下文信息进行知识外推

问题：英国现任首相是谁？



$$\begin{aligned} T(h, x) &:= P_1(h, x) \vee P_2(h, x) \\ P_1 &= \exists e, p \text{ Elec}(h, e) \wedge \text{Win}(e, p) \wedge \text{Leader}(p, x) \\ P_2 &= \exists c, p \\ &\quad \text{Parl}(h, c) \wedge \text{Largest}(c, p) \wedge \text{Leader}(p, x) \end{aligned}$$

- 路径1（选举）：英国→2024大选→获胜党派→党魁候选人
- 路径2（议会）：英国→下议院→最大党派→党魁候选人
- 两条链路均包含两个中间节点，并汇聚到同一候选人x

回答：英国现任首相是斯塔默

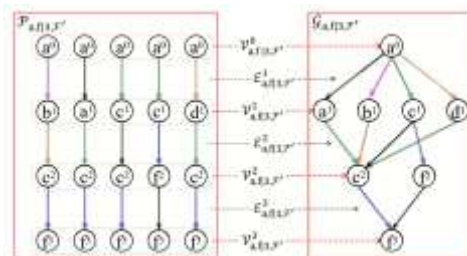
在新实体、新关系和分布变化下，将结构化上下文组织为可靠推理证据

结构化上下文作为推理证据 — 创新与成果

创新点：将关系路径与局部子图刻画为推理证据

证据建模：r-digraph 组织结构化关系证据

r-digraph: 将实体之间的多条关系路径组织为分层有向子图，在同一结构中保留关系顺序、共享中间实体与局部邻域结构。



r-digraph



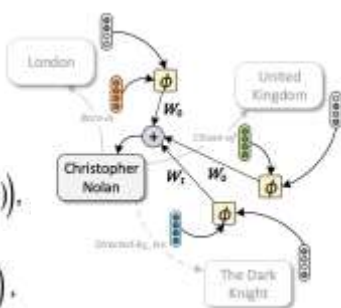
背景知识图谱

学习架构：从证据聚合到关系预测

检索与查询相关的关系路径和局部子图，利用 r-digraph 对多跳关系与邻域结构证据进行聚合建模，从而对候选实体完成可靠的外推关系预测。

$$h_{e_o}^{\ell}(e_q, r_q) = \delta \left(W^{\ell} \cdot \sum_{(e_s, r, e_o) \in \tilde{E}_{e_q}^{\ell}} \alpha_{e_s, r, e_o | r_q}^{\ell} (h_{e_s}^{\ell-1}(e_q, r_q) + h_r^{\ell}) \right),$$

$$\alpha_{e_s, r, e_o | r_q}^{\ell} = \sigma \left((w_{\alpha}^{\ell})^{\top} \text{ReLU} \left(W_{\alpha}^{\ell} \cdot (h_{e_s}^{\ell-1}(e_q, r_q) \oplus h_r^{\ell} \oplus h_{r_q}^{\ell}) \right) \right),$$



理论保证：关系规则的可学习性

QL-GNN

定理：给定查询 $(h, R, ?)$ ，用逻辑 $R(h, x)$ 描述规则，当且仅当 $R(h, x)$ 是 $CML[G, h]$ 中的逻辑时，QL-GNN可以学习该规则。

传统GNN

定理：给定查询 $(h, R, ?)$ ，用逻辑 $R(x, y)$ 描述规则，传统GNN可以学习的规则 $R(x, y)$ 可被拆分成关于 x 和 y 的一元逻辑公式的组合。

QL-GNN的能力强于一一般的链接预测方法

结构化上下文作为推理证据 — 效果与影响

效果：在新实体下实现稳定外推，并由关系子图和规则链路支撑解释

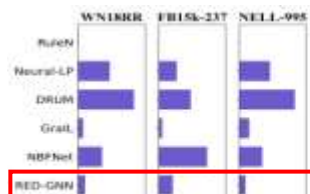
新实体预测效果

在新实体相关的预测任务上，RED-GNN可以实现稳定外推

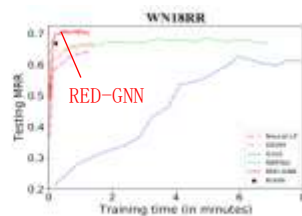
Performance on inductive reasoning with static KG.

		WN18RR				FB15k-237				NELL-995			
		V1	V2	V3	V4	V1	V2	V3	V4	V1	V2	V3	V4
MRR	RuleN	0.668	0.645	0.368	.624	.363	.433	.439	.418	0.615	0.385	0.381	.333
	Neural LP	0.649	0.635	0.361	0.617	0.325	0.389	0.400	0.396	0.610	0.361	0.367	0.261
	DRUM	0.666	0.646	0.380	0.627	0.333	0.395	0.402	0.410	.628	0.365	0.375	0.273
	GraIL	0.627	0.625	0.323	0.553	0.279	0.276	0.251	0.227	0.481	0.297	0.322	0.262
	CoMPLE	0.577	0.578	0.308	0.548	0.287	0.276	0.262	0.213	0.330	0.248	0.319	0.229
	NBFNet	.684	.652	.425	0.604	0.307	0.369	0.331	0.305	0.534	.410	.425	0.237
	InGram	0.277	0.236	0.230	0.118	0.293	0.274	0.233	0.214	0.584	0.358	0.308	0.221
	RED-GNN	.707	.686	.437	.634	.390	.472	.465	.465	.623	.496	.446	.373
	RuleN	63.5	61.1	34.7	58.3	30.9	34.7	34.5	31.8	54.5	30.4	30.3	24.8
	Neural LP	59.2	57.5	30.4	57.4	24.3	28.6	30.9	28.9	50.0	24.9	26.7	13.7
Hit@1	DRUM	61.3	59.5	33.0	58.6	24.7	28.4	30.8	30.9	50.0	27.1	26.2	16.3
	GraIL	55.4	54.2	27.8	44.3	20.5	20.2	16.5	14.3	42.5	19.9	22.4	15.3
	CoMPLE	47.3	48.5	25.8	47.3	20.8	17.8	16.6	13.4	10.5	15.6	22.6	15.9
	NBFNet	58.5	56.3	36.0	53.7	19.0	22.9	20.6	18.5	50.0	29.3	31.6	15.4
	InGram	13.0	11.2	11.6	4.1	16.7	16.3	14.0	11.4	50.0	25.3	19.9	12.4
	RED-GNN	64.6	62.9	38.5	58.9	32.4	37.2	36.9	37.1	52.5	39.8	35.0	26.0
	RuleN	63.5	61.1	34.7	58.3	30.9	34.7	34.5	31.8	54.5	30.4	30.3	24.8
	Neural LP	59.2	57.5	30.4	57.4	24.3	28.6	30.9	28.9	50.0	24.9	26.7	13.7
	DRUM	61.3	59.5	33.0	58.6	24.7	28.4	30.8	30.9	50.0	27.1	26.2	16.3
	GraIL	55.4	54.2	27.8	44.3	20.5	20.2	16.5	14.3	42.5	19.9	22.4	15.3

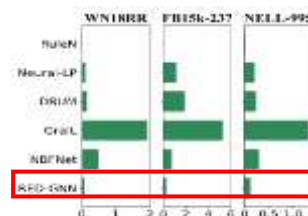
RED-GNN 模型参数少、训练效率高、推理速度快



(a) Model parameters (V1).



(b) Training curve (V1).



(c) Inference time on V1 (min).

可解释子图示例

RED-GNN方法的预测由子图和关键规则链路支撑。



种族 + 结婚地址 → 预测国籍



创作关联 + 系列关系 → 预测电影类型

药物安全应用 — 问题与需求

问题：风险**真实存在**，安全判断**要求高**，但**可用组合标签有限**

背景与需求

国际前沿：

- 420 亿美元/年成本；
- WHO 将用药安全列为全球患者安全挑战。
- Medication Without Harm全球挑战：目标减少 50% 严重、可避免伤害。

国内需求：

- 患者安全专项行动（2023-2025）持续强调高警示药品管理。
- 药物警戒质量管理规范：强调风险识别、评估与监测。



典型高风险药物组合

药物组合	可能严重后果
头孢曲松 + 葡萄糖酸钙	形成头孢曲松-钙盐沉淀
呋塞米 + 阿米卡星	增加耳毒性与肾毒性风险
阿曲库铵 + 阿米卡星	加重或延长呼吸抑制
奥美拉唑 + 氯吡格雷	降低氯吡格雷疗效
阿司匹林 + 氯吡格雷	增强血小板抑制效应

五种主要药物相互作用组合及其后果



组合空间巨大

药物两两组合
数量快速增长



风险低频高危

严重不良反应
样本少、覆盖难



临床验证

成本高
风险依赖剂量
病人状态与合并用药



监督缺口明显

联用标签有限
训练仍需大量
监督

真实风险空间 >> 已知标注空间

用药安全风险客观存在，但真实风险空间远大于可用组合标签空间

药物安全应用 — 方案思路

方案思路：将分散医学证据组织为结构化上下文

多源医学证据

药物/靶点/疾病
副作用/文献/病例

结构化上下文构建

组织实体关系
形成结构上下文

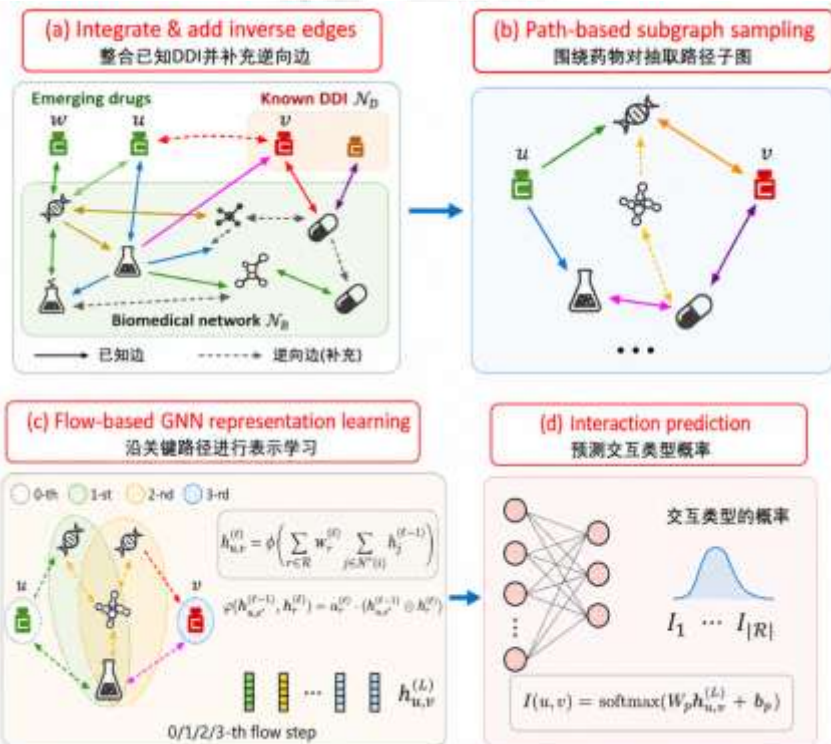
关键证据抽取

检索相关证据
识别关键路径

风险预测与证据输出

判断相互作用类型
输出风险依据

方法—EmerGNN



结构化上下文为什么有效

共性规律

共享基因/副作用
/化合物

用药规则

多跳关系路径
体现机制依据

降低数据
依赖

利用共享实
体关系补充
监督信号

识别关键
路径

从候选路径
中筛出关键
作用链条

支撑
复核

形成可追踪的
判断依据支持
专家复核

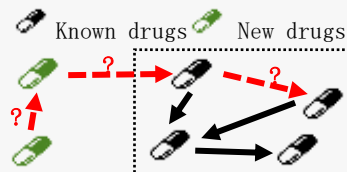
针对数据稀缺的新药相互作用的轻量级且可解释的模型

药物安全应用 — DDI Agent

方案思路：从固定证据路径，升级为按需自适应路由证据、专家与推理流程

DDI 预测—经典小样本问题

- 核心定义与风险：药物间相互影响药效，可致严重不良反应，致死率0.32%。
- 现实挑战：数据稀疏——标注成本高、验证昂贵，导致可用标签极少，新药上市无历史记录



知识来源

分子结构

SMILES 及结构线索

生物学关系

靶点、已有DDI数据

语义上下文

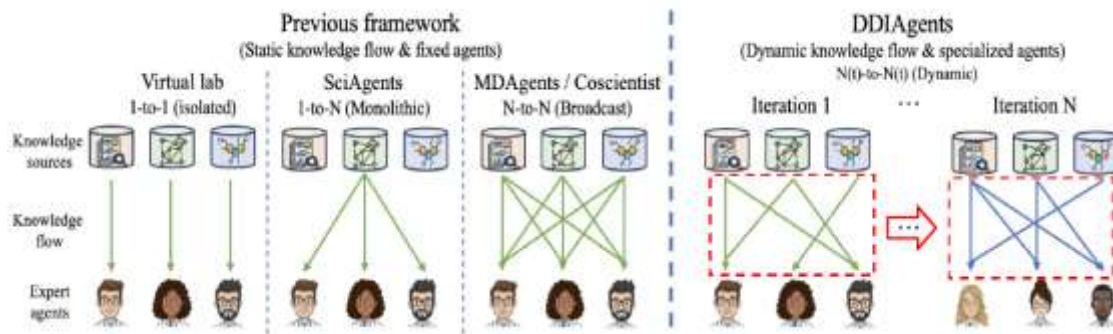
LLM生成的机制和DDI类型摘要

路由规则

- 根据专家能力将证据分配给专家
- 将专家间的分歧与改进用作监督信号的替代
- 将来源到专家的连接视为变量，而非固定管道

为什么固定知识流不够？

- 分子结构方法：强于化学信号，弱于跨来源证据整合
- LLM 方法：可利用文本知识，但证据选择缺乏控制
- 多智能体方法：可接入多源知识，但知识流通常固定



将source-to-expert边视为变量，让多源医学证据按药物对与迭代状态动态流动

药物安全应用 — DDIAgents框架与验证

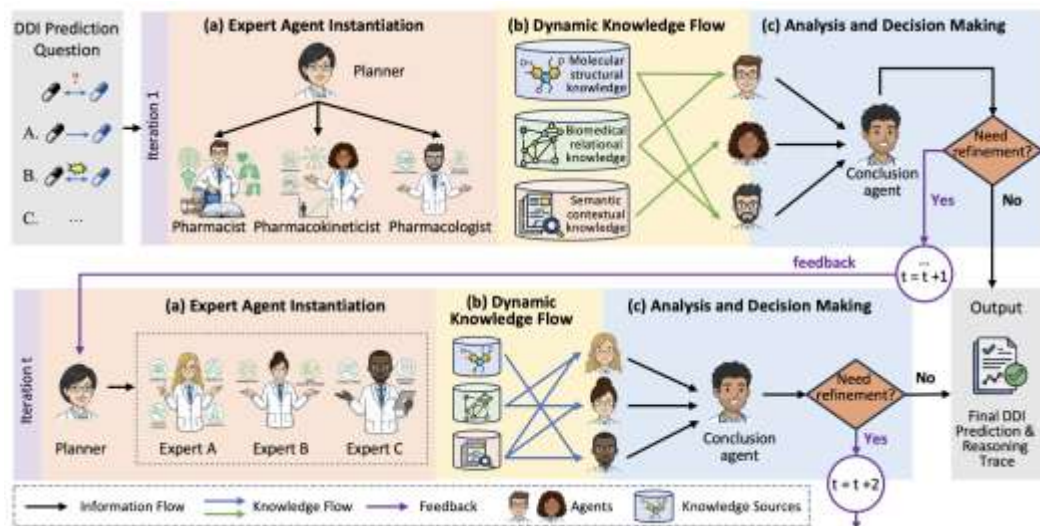
用专家实例化、动态知识流和两轮修正，弥补新药 DDI 场景下的标签稀缺

1 专家智能体实例化 2 动态知识流 3 分歧修正

- Planner 根据当前药物对定制化创建专家团队

- 分子结构/生物医学关系/语义机制按药物对的需求自适应路由

- 综合各专家分析结果；专家冲突触发critique与第二轮推理



药物对 → 专家分析 → 分歧检测 → 更新知识流 → 最终 DDI 判断
优势：每例自适应调整，无需额外实验室数据收集。

基准实验与突出表现

SOTA
真实 DDI 基
准全面领先

+18.5%
小样本跨机
制外推提升

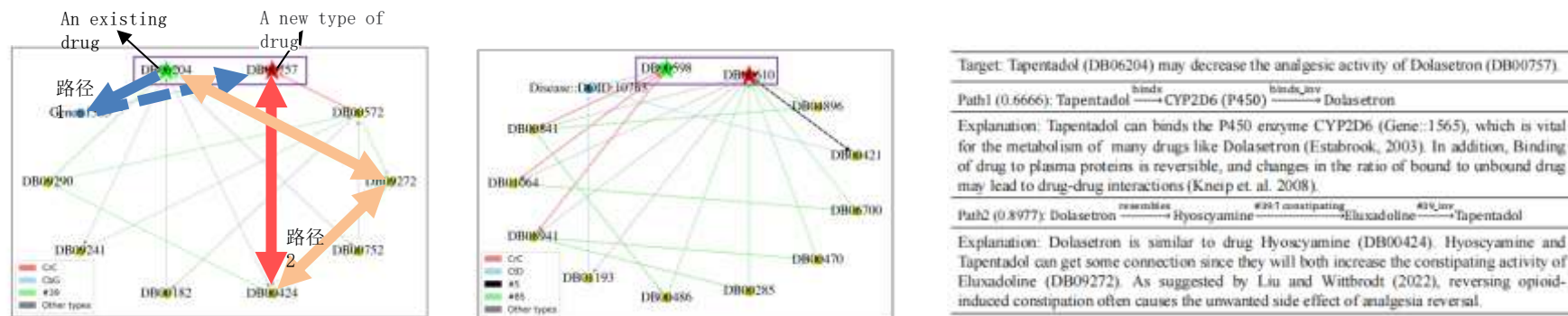
100%
专家级推导
轨迹可解释

- ▶ **全面超越传统基准：**显著优于基于特征 (Feature)、传统图神经网络 (GNN)、大模型直接预测 (LLM Direct) 及单 Agent 架构。
- ▶ **消除伪证据干扰：**通过机制筛选机制，将交互概率与医学机理解耦，有效避免大模型幻觉引起的伪证据链。
- ▶ **专家复核友情支持：**自动输出“结构特征-靶点结合-临床文献”多级推理轨迹，支持专家复查与反事实验证。

小样本药物安全预测，不只是更多标签，是可路由、可解释的结构化上下文

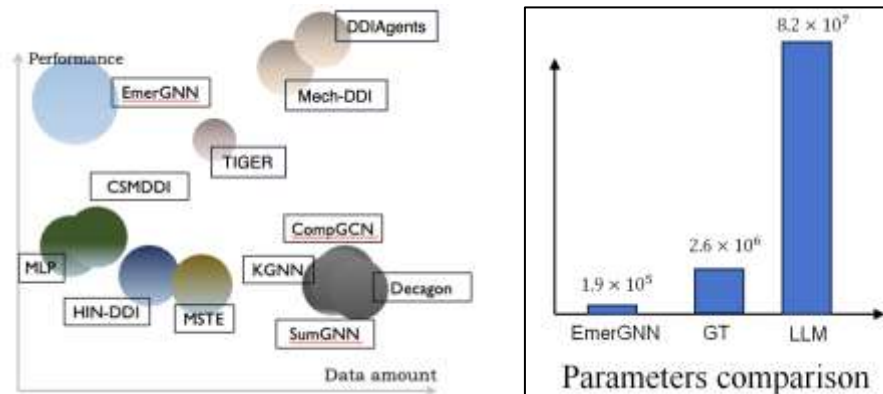
药物安全应用 — 效果与价值

价值：面向数据稀缺的新药相互作用，形成轻量、可解释的预测方法



结构化路径支持新药相互作用预测，并为专家复核提供可追踪依据

Datasets		DrugBank			TWO SIDES		
Type	Methods	F1-Score	Accuracy	Kappa	PR-AUC	ROC-AUC	Accuracy
DF	MLP	21.1±0.8	46.6±2.1	33.4±2.5	81.5±1.5	81.2±1.9	76.0±2.1
	CSMDDI	45.5±1.0	62.6±2.8	55.0±3.2	73.2±2.6	74.2±2.9	69.9±2.2
GF	HIN-DDI	37.3±2.9	58.9±1.4	47.6±1.8	81.9±0.6	83.8±0.9	79.3±1.1
Emb	MSTE	7.0±0.7	51.4±1.8	37.4±2.2	64.1±1.1	62.3±1.1	58.7±0.7
	KG-DDI	26.1±0.9	46.7±1.9	35.2±2.5	79.1±0.9	77.7±1.0	60.2±2.2
GNN	CompGCN	26.8±2.2	48.7±3.0	37.6±2.8	80.3±3.2	79.4±4.0	71.4±3.1
	Decagon	24.3±4.5	47.4±4.9	35.8±5.9	79.0±2.0	78.5±2.3	69.7±2.4
	SumGNN	35.0±4.3	48.8±8.2	41.1±4.7	80.3±1.1	81.4±1.0	73.0±1.4
	EmerGNN	62.0±2.0	68.6±3.7	62.4±4.3	90.6±0.7	91.5±1.0	84.6±0.7
GT	TIGER	60.5±1.8	67.4±3.2	61.0±3.9	88.9±0.8	90.2±1.1	83.2±0.9
LLM	Mech-DDI	63.4±1.7	69.8±3.1	63.6±3.8	91.2±0.6	92.1±0.8	85.5±0.7
Agg	DDIAgents	65.1±1.5	71.2±2.9	65.0±3.6	92.4±0.5	93.2±0.7	86.8±0.6



组织多源医学证据，降低训练数据依赖并提升结果可解释性

从模型泛化走向智能体适应

结构化上下文从静态关系拓展到动态交互过程，连接少量样本下的模型泛化与少量经验下的智能体适应。

已有基础： 模型泛化

快速泛化

少量样本下，模型为何能借助上下文快速泛化？

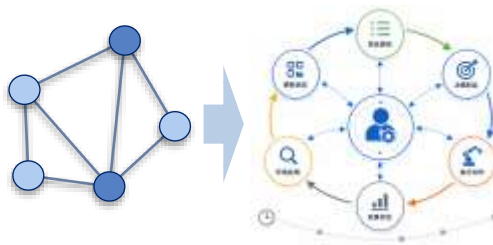
高效适配

如何将上下文转化为可计算、可搜索、可迁移的适配能力？

可靠推理

如何在新对象、新关系与分布变化下进行可解释推理？

关键变化： 从静态关系到动态经验



静态样本关系
(数据层)

动态交互经验
(状态、轨迹与反馈)

显式监督更稀疏，过程信息更丰富，结构化上下文更加关键

核心转化：静态关系 → 过程轨迹
→ 可复用结构化经验

未来方向： 智能体适应

机理深化：解释模型泛化

研究上下文在模型内部的表示、利用机制及其对训练稳定性的影响。

场景拓展：支持智能体适应

将结构化上下文拓展到交互过程，支撑长期决策、经验复用与可靠执行。

核心转化：静态关系 → 过程轨迹 → 可复用结构化经验

近期方向一：系统层动态适应（MASF1y）

从固定多智能体系统走向测试时可重构的系统组织

任务差异与执行状态变化不断出现，如何在不断训练的情况下让多智能体系统持续适应？

为什么需要？

- 不同任务需要不同角色与协作流程
- 执行中会出现工具失败与协调错误
- 部署后重新标注和训练成本高

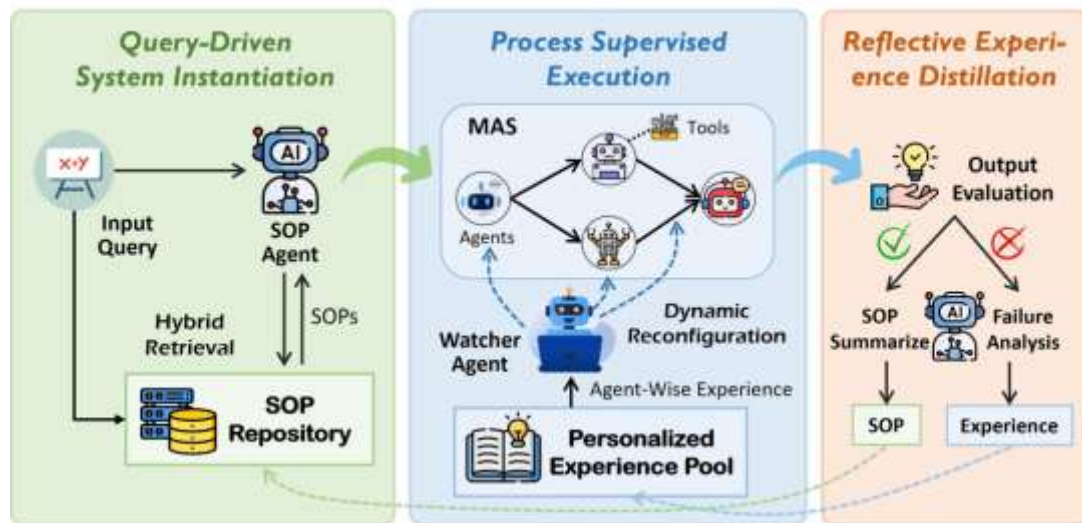
结构化上下文

- SOP 模板：保存成功协作模式
- PEP 经验池：保存失败模式与修正策略
- Watcher：根据执行状态触发纠错或替换

无需参数更新

执行时纠错

经验可复用



任务请求 → SOP 检索 → 任务特定 MAS → Watcher 监督纠错 → 经验蒸馏

MASF1y 将结构化上下文从“数据样本关系”扩展为“系统组织与执行经验”

近期方向二：模型层自我改进（SGALM）

从外部标注走向模型内部的自生成—自判别监督

核心问题：当高质量对齐数据稀缺时，模型如何生成并筛选自己的训练信号？

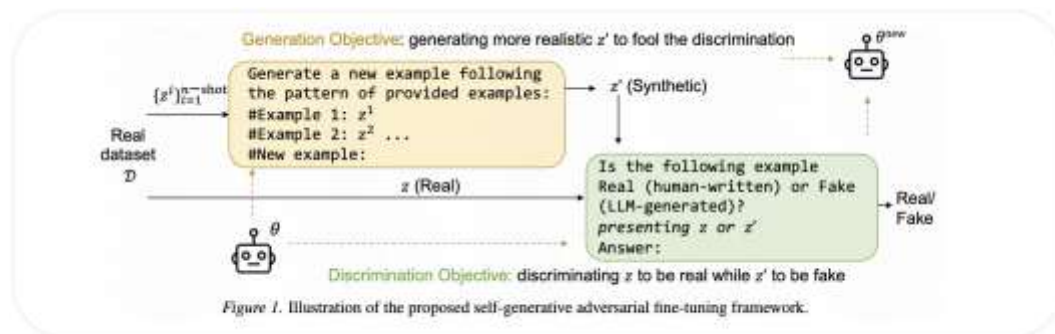
自我改进的风险

- 高质量 alignment 数据稀缺且昂贵
- 自生成样本可能偏离真实任务分布
- self-evaluation bias 会放大模型已有错误

结构化上下文

- 少量真实样本提供分布锚点
- 合成候选数据扩展监督信号
- 内部判别概率过滤低质量样本

同一 LLM 同时生成、判别，并从二者中学习。



少量真实样本 → Generator 生成候选数据 → Discriminator 判断真伪 → 交替更新

无需外部 Judge

自生成监督

内部质量控制

SGALM 将结构化上下文从“外部证据”扩展为“模型内部的生成—判别信号”

报告小节

❖ 时代问题

- ❖ AI 从静态预测走向智能体执行，小样本学习从少量样本扩展到少量经验。

❖ 核心观点

- ❖ 有限监督不等于有限信息；属性、关系、过程、约束与交互历史可以构成结构化上下文。

❖ 研究路径

- ❖ 结构化上下文作为隐式监督、适配信号与推理证据，支持快速泛化、高效适配和可靠外推。

❖ 未来方向

- ❖ 从静态样本关系走向动态交互经验，支撑智能体持续学习、可靠推理与自主适应。

感谢各位老师，请批评指正！