



清华大学 交叉信息研究院

Institute for Interdisciplinary Information Sciences, Tsinghua University

学术科研简报

IIS Academic Newsletter

2026.01-06



目 录

人工智能领域 (04-35)

人工智能大模型	04
具身智能和机器人	17
计算机图形学	30
人工智能安全	31
人工智能理论	34

计算机科学领域 (38-53)

计算机系统结构	38
数据库管理系统	43
计算经济学	46
计算机安全	47
密码学	48
算法理论	53

量子信息领域 (55-70)

量子模拟、量子传感	55
量子信息、机器学习	60
离子阱量子网络	68
量子计算、凝聚态物理	69

人工智能



一、人工智能大模型

主要完成人：李建研究组、马恺声研究组、陈建宇研究组、房智轩研究组、贺天行研究组、张景昭研究组

AdaRoPE: 不同注意力头应有不同的旋转频率与缩放因子

旋转位置编码（Rotary Position Embedding, RoPE）是当前大语言模型中广泛采用的位置编码方案，被 LLaMA、Qwen、Gemma 等主流模型所使用。标准 RoPE 通过对每个 token 的查询和键向量施加旋转矩阵来编码位置信息。然而，现有的 RoPE 实现对所有注意力头使用统一的旋转频率和缩放因子，存在两个关键局限：第一，统一的频率分配导致嵌入维度利用不充分——不同功能的注意力头（如语义头仅利用低频、偏移头利用高频）被迫共享同一频率表，造成注意力表征能力的浪费；第二，在长上下文扩展中，YaRN 等方法对所有头统一缩放频率和注意力温度，忽略了不同头类型的异质性需求。

针对上述局限，李建团队提出了 AdaRoPE 方法，为每个注意力头配备可学习的旋转频率（AdaFreq）和长度感知的注意力缩放因子（AdaScale）。AdaFreq 为每个头的每个维度学习独立的旋转角速度，使不同头能够自适应地选择最适合其功能的频率范围。AdaScale 则为每个头引入与序列长度相关的注意力温度缩放，使检索型注意力头能保持尖锐的注意力分布以精确定位目标 token，而全局聚合型注意力头则能随上下文增长而展平注意力分布。理论分析表明，不同窗口宽度的注意力头确实需要不同的频率范围和缩放因子。AdaRoPE 作为即插即用的轻量级扩展，可直接替换标准 RoPE，与 FlashAttention 和标准 KV 缓存完全兼容，额外参数量仅占模型总参数的十万分之一级别。

在预训练实验中，团队在 430M 至 2.7B 参数规模的多种模型架构（LLaMA、OLMoE、Qwen3 Gated Attention）上验证了 AdaRoPE 的有效性。结果表明，AdaRoPE 在七项自然语言理解基准上的平均准确率比标准 RoPE 高出 0.91 个百分点，且在所有模型架构上均取得最优性能。在上下文扩展实验中，团队将 Llama-3-8B 和 SmolLM-1.7B 的上下文窗口从 8k 扩展至 64k。在零样本外推设定下，仅微调频率和缩放参数（冻结主干网络），AdaRoPE 在 LLaMA-8B 的 64k 长度检索任务上达到 51.87 的准确率，而 YaRN 仅为 12.80。在长上下文持续预训练设定下，采用两阶段优化策略（先预热 AdaRoPE 参数，再联合训练 LoRA 与 AdaRoPE），AdaRoPE 在 LLaMA-8B 上达到 69.09 的 64k 检索准确率，显著优于 YaRN 的 61.60。

该成果研究论文：Shaowen Wang, Yuke Zheng, Tansheng Zhu, Shuang Chen, Shaofan Liu, Suncong Zheng, and Jian Li,

“AdaRoPE: Not All Attention Heads Should Rotate and Scale Equally,” ICML 2026.

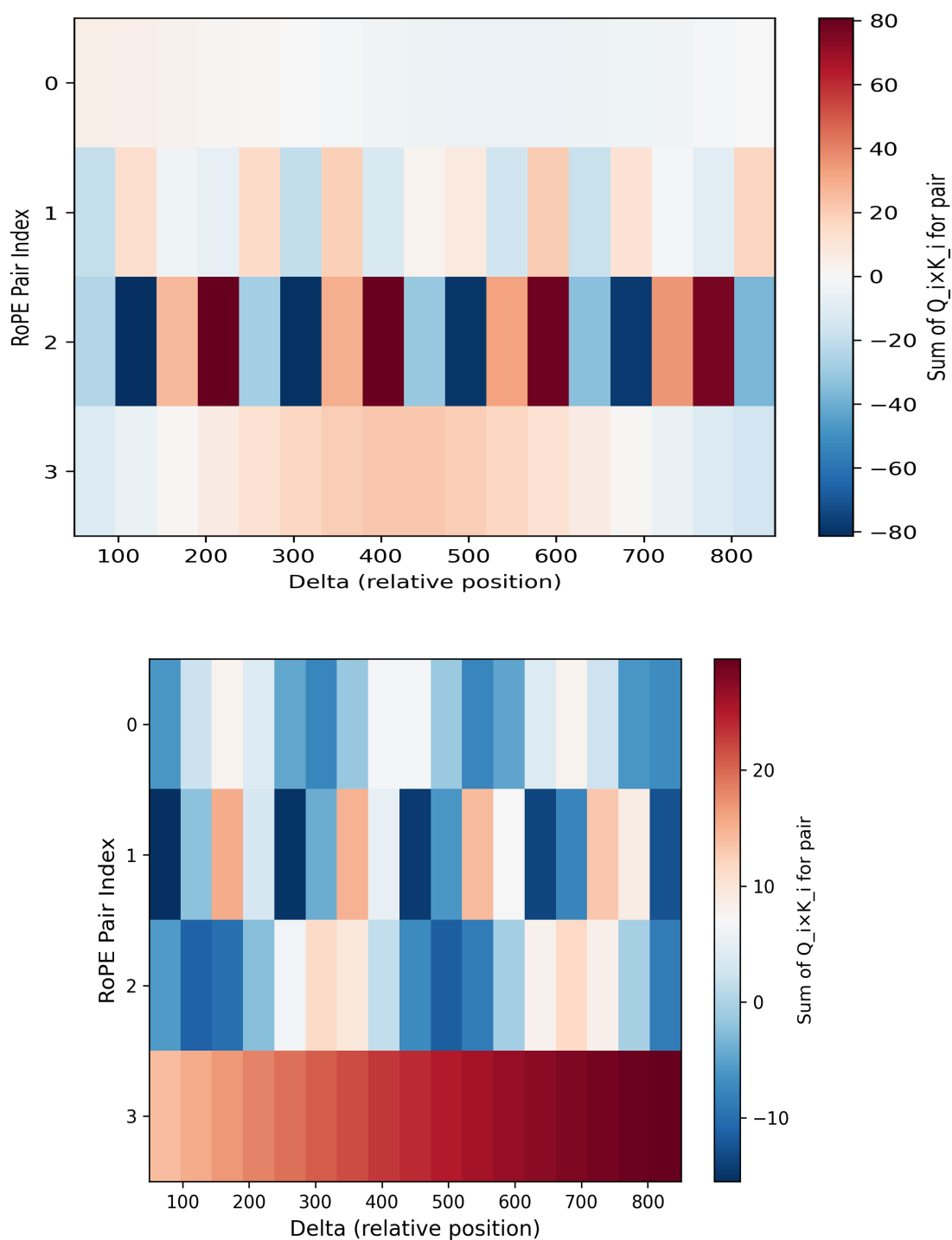


图 1: RoPE 各维度对注意力 logit 的贡献热力图。左图为标准 RoPE, 仅有少数维度被有效利用, 呈现高度稀疏的频率使用模式; 右图为 AdaRoPE (可学习频率), 通过自适应调整各头的旋转频率, 显著提升了所有注意力维度的利用率。

ES-dLLM：基于早跳过的扩散大语言模型高效推理

扩散大语言模型（diffusion large language models, dLLMs）近年来成为自回归大语言模型之外的重要技术路线。与逐词生成的自回归模型不同，dLLM 通过多轮去噪逐步将掩码位置还原为文本，能够利用双向上下文并具备并行生成潜力。然而，这类模型在每一轮推理中通常需要处理完整输入序列，导致大量重复计算，成为制约开源 dLLM 推理效率的关键瓶颈。

马恺声团队围绕扩散大语言模型的推理冗余展开研究，并提出无需额外训练的推理加速框架 ES-dLLM。该方法首先分析 dLLM 生成过程中的置信度变化和隐藏状态变化，发现多数位置在连续迭代中的变化接近于零。基于这一观察，ES-dLLM 综合利用上一轮置信度和当前中间张量变化估计 token 重要性，在 Transformer 早期层跳过低重要性位置的计算，并通过局部缓存更新并复用未更新位置的 KV 缓存，从而减少每轮推理中的冗余计算。

团队在 LLaDA-8B 和 Dream-7B 两个代表性扩散大语言模型上进行了系统验证。实验显示，ES-dLLM 在 NVIDIA H200 GPU 上最高可达到 226.57 和 308.51 tokens/s，相比原始实现获得 5.6 倍至 16.8 倍加速，相比当前先进 KV 缓存方法 DualCache 最高仍有 1.85 倍提升，同时基本保持在多个数据集上的生成质量。进一步实验表明，该方法还可与并行解码、稀疏注意力等现有加速技术结合，为扩散大语言模型的高效部署提供了通用组件。

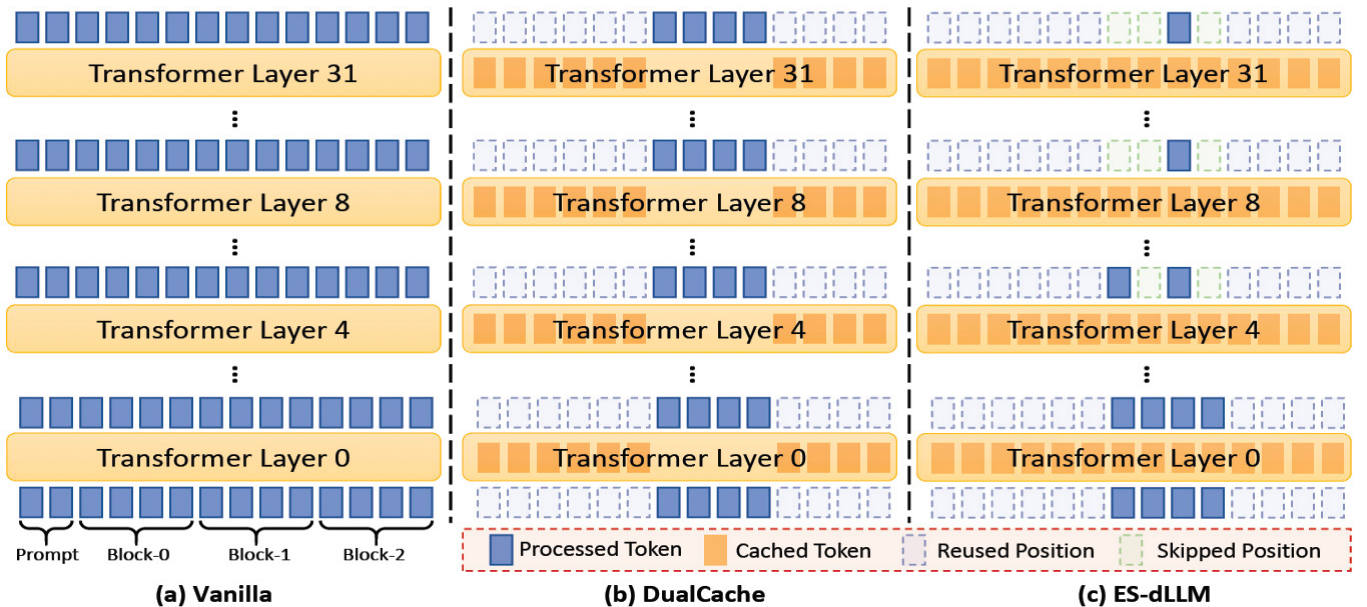


图 1：ES-dLLM 早跳过推理加速机制示意以及与现有方法的对比

该成果研究论文：Zijian Zhu, Fei Ren, Zhanhong Tan, and Kaisheng Ma, “ES-dLLM: Efficient Inference for Diffusion Large Language Models by Early-Skipping,” ICLR 2026.

VLA 模型和世界模型的迭代优化

视觉 - 语言 - 动作模型（Vision-Language-Action, VLA）已经在机器人操作任务中展现出较强的泛化能力，但其性能进一步提升往往依赖真实世界中的在线交互数据。然而，在真实机器人环境中收集策略执行轨迹成本很高，需要人工重置场景、监控机器人执行并处理失败情况，难以大规模扩展。

陈建宇团队研究了视觉 - 语言 - 动作策略与世界模型的迭代协同改进问题，提出了 VLAW（Iterative Co-Improvement of Vision-Language-Action Policy and World Model）框架。该方法首先让 VLA 策略在真实世界中进行少量在线执行，收集包含成功与失败情况的真实轨迹；随后利用这些轨迹对动作条件世界模型进行后训练，使其更好地刻画目标任务中的物理动力学和失败模式。经过真实交互数据校准后的世界模型再被用于生成大规模合成执行轨迹，并通过微调后的视觉 - 语言奖励模型自动筛选成功轨迹，最终将这些真实与合成的成功数据用于继续训练 VLA 策略。

团队在 DROID 真实机器人平台上验证了 VLAW 的有效性，实验覆盖堆叠积木、翻开书本、擦除白板痕迹、舀取零食和画圆等多类接触丰富或涉及可变形物体的操作任务。实验表明，使用在线策略轨迹微调后的世界模型不仅显著提升了视频预测质量，还能更准确地区分真实操作中的成功与失败结果，减少仅依赖专家示范时产生的过度乐观预测。进一步地，VLAW 在多任务真实机器人实验中相较基础 VLA 策略取得了 39.2% 的绝对成功率提升，其中由世界模型生成的合成轨迹带来了 11.6% 的额外性能提升。该研究表明，随着视频生成模型和机器人交互数据的发展，基于世界模型的想象式训练有望成为提升通用机器人策略性能的重要范式。零食和画圆等多类接触丰富或涉及可变形物体的操作任务。实验表明，使用在线策略轨迹微调后的世界模型不仅显著提升了视频预测质量，还能更准确地区分真实操作中的成功与失败结果，减少仅依赖专家示范时产生的过度乐观预测。进一步地，VLAW 在多任务真实机器人实验中相较基础 VLA 策略取得了 39.2% 的绝对成功率提升，其中由世界模型生成的合成轨迹带来了 11.6% 的额外性能提升。该研究表明，随着视频生成模型和机器人交互数据的发展，基于世界模型的想象式训练有望成为提升通用机器人策略性能的重要范式。

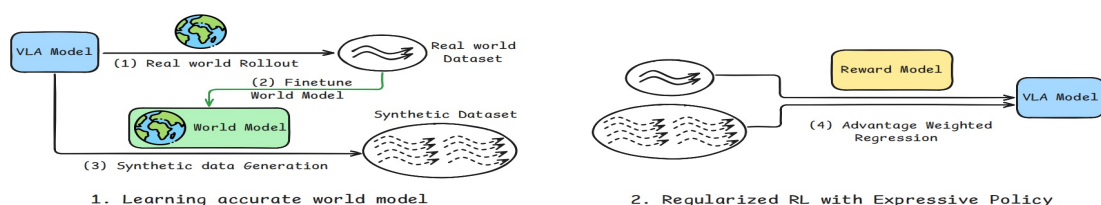


图 1：VLA 世界模型迭代示意图

该成果研究论文：Yanjiang Guo, Tony Lee, Lucy Xiaoyang Shi, Jianyu Chen, Percy Liang, Chelsea Finn VLAW: Iterative Co-Improvement of Vision-Language-Action Policy and World Model, ICML 2026.

基于选择性延迟路由的端侧—云端语言模型低成本高性能协作

房智轩团队针对开源大语言模型（LLM）推理服务在成本与性能之间难以兼顾的问题，提出了选择性延迟路由，以实现端侧与云端模型之间的低成本、高性能协作。目前，大语言模型服务主要通过两种方式获取：（i）付费访问云端托管的大语言模型，这类模型能力强大但会带来不可忽视的成本；（ii）在个人设备或小型集群上部署小语言模型，这类模型能力相对较弱，但足以处理较为简单的任务。

为此，文章作者提出了选择性延迟路由这一范式，从而在成本与性能之间实现平衡的折衷。在该框架中，用户请求首先由本地小语言模型处理，它不仅生成初步响应，还提供请求的丰富语义表示。随后，一个轻量级决策模块利用这些信息，决定是采纳初始响应，还是将请求路由到最合适的远程大语言模型以获取更高质量的响应。针对多种模型架构与家族（涵盖小语言模型和大语言模型）以及跨多个任务场景的数据集所进行的大量实验表明，研究人员的方法在广泛的成本 - 性能权衡指标下，始终优于现有的多模型协作方法。

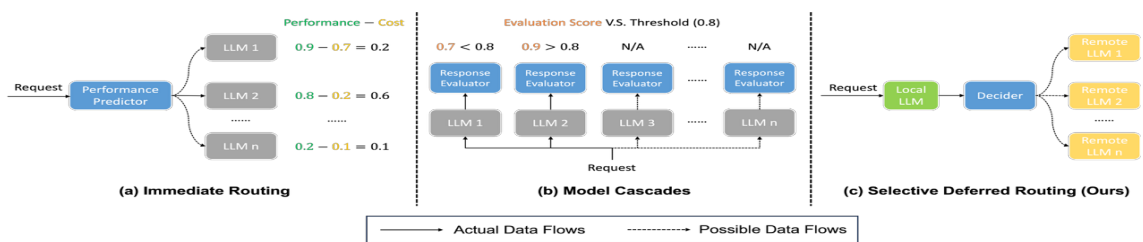


图 . 选择性延迟路由（该文提出范式）与现有的两种模型协作范式的对比

该成果研究论文：Qijun Miao and Zhixuan Fang, “Selective Deferred Routing: Enabling Cost-Efficient Collaboration between Local SLMs and Remote LLMs” ,ICML 2026.

对抗数据投毒攻击的近乎最优鲁棒联邦学习

房智轩团队提出了一种面向联邦学习数据投毒攻击的近乎最优鲁棒防御机制。联邦学习允许多个节点在不共享原始数据的情况下协同训练模型，但若部分节点的本地数据被恶意篡改，例如标签翻转或后门植入，最终模型可能在看似正常的训练流程中被污染。

论文聚焦“节点多、单节点数据少”的现实场景：单个节点难以判断自身数据是否异常，但正常节点整体仍包含足够统计信息。团队提出协同检测机制，先训练辅助判别模型刻画节点间数据分布差异，并据此分配可信权重，再利用加权数据训练目标模型。相比依赖高维梯度异常检测的传统鲁棒聚合方法，该机制更直接利用了数据投毒攻击的结构特征。

理论上，论文给出了数据投毒攻击下模型性能损失的下界，并证明算法在 IID 与非 IID 设置下均可渐近匹配该下界。同时，论文提出了有效投毒率（EPR）的新指标，用以衡量攻击者除降低目标模型准确性之外，对特定数据预测结果的扭转能力。实验表明，该方法在标签翻转和后门攻击中，相比 geometric median、iterative filtering、Krum 等已有方法，可保持更高测试准确率并降低攻击成功率。

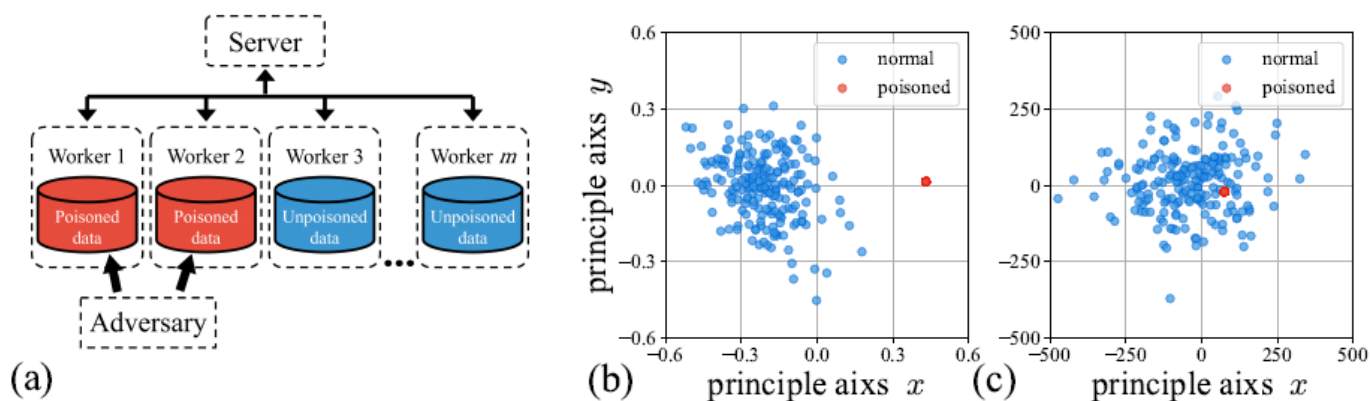


图. (a) 联邦学习系统中的数据投毒攻击。(b-c) 同一数据集在不同神经网络上计算得到的梯度主成分轴

该成果研究论文: Jingfan Yu and Zhixuan Fang, “Near Optimal Robust Federated Learning Against Data Poisoning Attack”, ICLR 2026.

安全约束下线性多臂老虎机中遗憾与统计推断的帕累托最优权衡

房智轩团队针对在线决策中累积遗憾、统计推断与安全约束三者之间长期被忽视的内在张力，提出了首个在三项目标上同时达到帕累托最优的线性随机老虎机算法 SERMiSC。在临床试验、移动健康、算法交易等高风险在线决策场景中，传统方法只关注累积回报，却忽视了对未知参数 θ^* 的精确估计——而后者不仅直接决定了平均处理效应 (ATE) 等核心统计量的可靠性，更决定了很多实验结论能否外推到更广人群与未来场景，且任何收益若以违反安全约束为代价都不可接受。然而，现有工作几乎从未在统一框架下刻画“安全—遗憾—推断”三者的本质权衡。为此，团队从信息论角度推导出统计估计误差在给定遗憾与安全预算下的极小极大下界，首次刻画了三者的帕累托前沿；进而设计两阶段算法 SERMiSC：先以代价感知的 Gibbs 策略安全地获得高精度参数估计，再切换至基于虚拟队列的 pessimistic-optimistic 策略最小化遗憾。理论分析证明 SERMiSC 精确匹配该下界、达到帕累托最优；仿真实验进一步表明：在与 OFUL、PO 等代表性算法相当的遗憾与安全表现下，SERMiSC 显著降低了对未知参数的估计误差，为安全敏感的在线决策提供了具有理论保证的新工具。



图. 累积遗憾、安全违反与统计估计误差三者相互掣肘、构成一个权衡，而 SERMiSC 通过两阶段设计精确匹配其理论下界，达到帕累托最优

该成果研究论文：Yuming Shao and Zhixuan Fang, “The Pareto-optimal Trade-off between Regret and Statistical Inference in Linear Stochastic Bandits under Safety Constraints”, ICML 2026.

大语言模型密码学能力综合评测

现代密码学是网络安全、隐私保护和数字基础设施中的核心技术，覆盖从基础数学理论到实际系统实现的多层次能力。随着大语言模型在数学推理、代码生成和自动化问题求解方面快速发展，一个重要问题逐渐凸显：现有大语言模型是否真正具备处理密码学问题的能力。该问题不仅关系到人工智能辅助密码学研究的可行性，也关系到未来网络安全场景中自动化漏洞分析、密码协议审查和安全评估工具的可靠性。

贺天行团队研究了大语言模型在密码学任务中的能力边界，并构建了综合性评测基准 AICrypto。该基准覆盖选择题、网络安全竞赛题和证明题三类任务，分别用于评估模型对密码学基础概念的掌握、对实际密码实现漏洞的利用能力，以及对形式化密码学论证的推理能力。为提高评测的严谨性，团队组织密码学专家对题目进行人工筛选、改写、验证，并为证明题设计详细评分标准和参考解答，以支持自动化评分和人类专家基线对比。

团队基于 AICrypto 评测了 17 个主流大语言模型，并与人类密码学专家表现进行比较。实验结果表明，当前最先进的大语言模型已经能够在基础概念记忆、常见密码漏洞利用和常规证明任务上接近甚至超过人类专家，但是对数学概念和实际攻击逻辑缺乏深入的理解。AICrypto 为系统评估人工智能的密码学能力提供了一个较为完整的基准，也为未来发展面向密码学研究、密码工程审查和安全证明分析的大语言模型工具提供了重要参考。

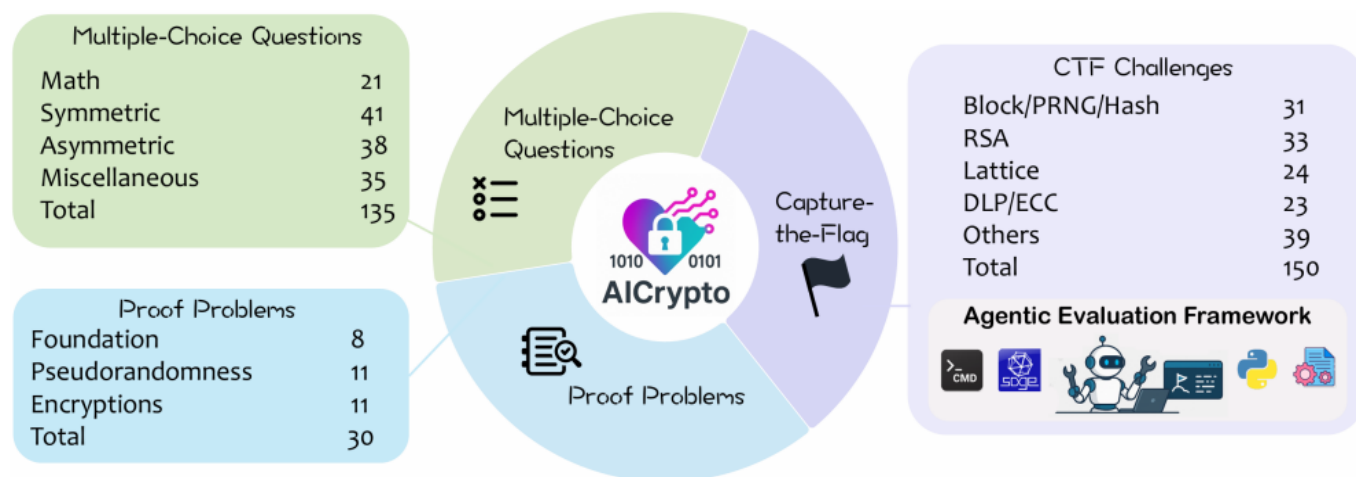


图 1： AICrypto benchmark 示意图

该成果研究论文：YuWang, Yijian Liu, Liheng Ji, Han Luo, Wenjie Li, Xiaofei Zhou, Chiyun Feng, Puji Wang, Yuhan Cao, Geyuan Zhang, Xiaojian Li, Rongwu Xu, Yilei Chen, Tianxing He “AICrypto: Evaluating Cryptography Capabilities of Large Language Models.” ICML 2026.

语言模型下的模型编辑逆向工程

模型编辑作为一种免训练的大语言模型参数更新方案，在事实更新，偏好对齐，信息擦除等领域中应用广泛。然而该技术潜在的数据泄露风险此前尚未得到充分探讨，针对这一空白，贺天行团队研究了主流的“定位后编辑”范式，首次证实了编辑更新的低秩结构可充当一种侧信道，从而实现对被编辑数据的逆向恢复。

团队首先从理论上证明了参数更新矩阵的行空间编码了被编辑主语独特的键向量，随后针对性设计了“键空间重构 - 熵值衰减 Key Space ReconsTruction-then-Entropy Reduction (KSTER)”的两阶段攻击框架，使攻击者在无需知晓原始提示词的情况下精准恢复出被编辑的数据。该攻击通过在第一阶段对参数更新矩阵进行谱分析，重构出键空间的线性子空间以精确锁定目标主语；随后在第二阶段，利用编辑前后模型预测置信度的剧烈变化（即信息熵的显著缩减），成功重构出与之关联的语义上下文。

为了缓解这一泄露风险，团队还提出了“子空间伪装 Subspace Camouflage”的主动防御策略；该策略通过在更新计算中注入由无关主语构成的“语义诱饵”来混淆键空间，从而实现在不降低编辑效用的同时有效误导攻击者。

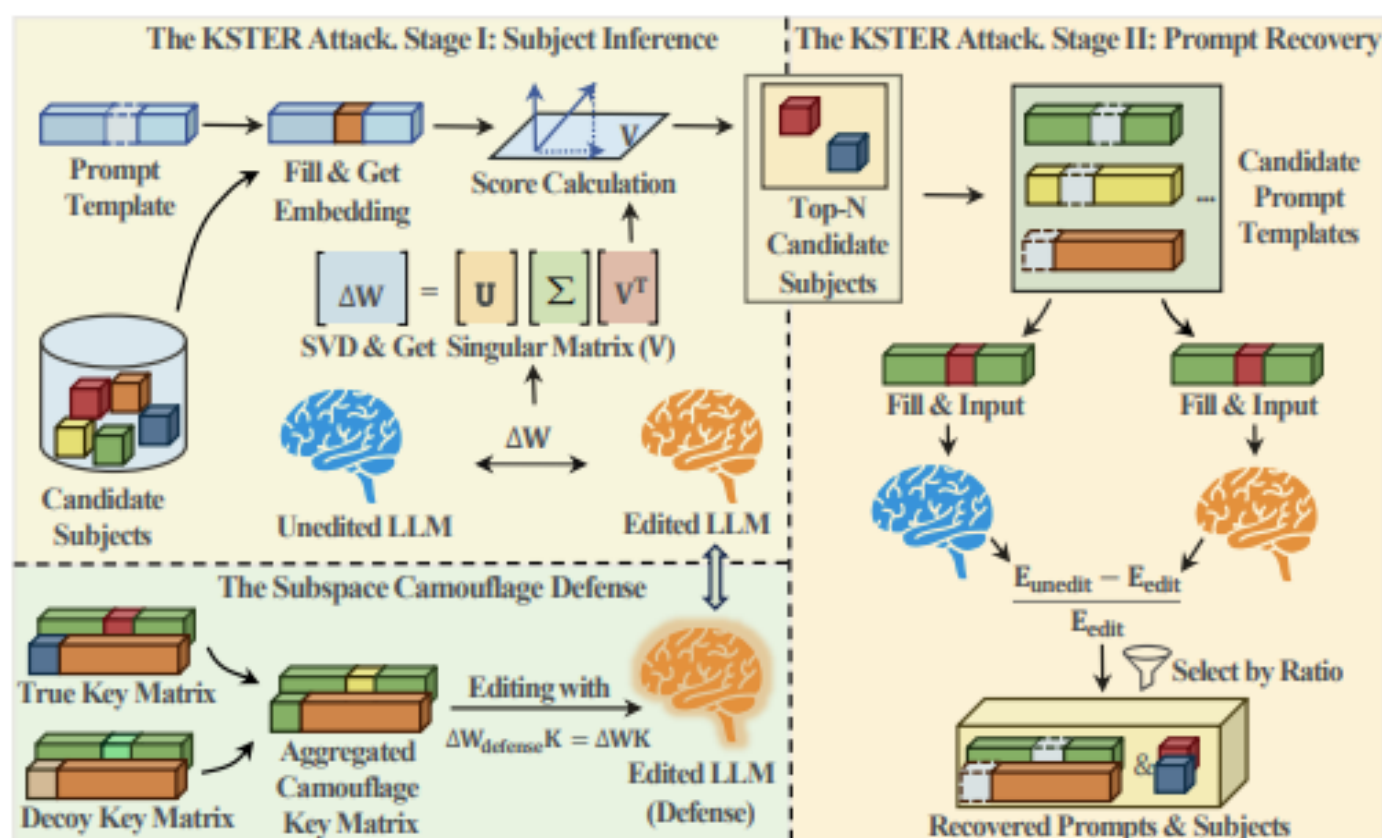


图 1：收敛动力学示意图

该成果研究论文：Zhiyu Sun, Minrui Luo, Yu Wang, Zhili Chen, Tianxing He, “Reverse-Engineering Model Editing on Language Models,” ICML, July 2026.

面向大语言模型训练的数据配比优化

大语言模型通常需要混合通用知识、代码、数学、中文、英文等多类数据进行训练，不同数据源的配比会显著影响模型在推理、数学、编程和知识问答等任务上的最终表现。但直接在目标大模型上搜索最优数据配比代价很高，而只在小模型上得到的最优配比又不一定能可靠迁移到更大模型。这给大模型训练中的数据选择带来了很大困难，因为研究者需要在有限计算资源下，同时控制数据配比搜索成本、模型规模外推误差和下游任务表现。

张景昭团队研究了大语言模型中训练阶段的数据配比优化问题，使研究者在不需要对目标大模型进行大量完整训练的情况下，能够通过设计的“容量感知数据配比定律” Capacity-Aware Mixture Law (CAMEL)，预测不同模型规模下各类数据配比对验证损失和下游能力的影响，从而高效选择适合目标模型规模的数据混合方案。该方法将数据配比优化理解为模型容量在不同内在能力领域之间的分配问题，显式建模模型规模与数据配比之间的非线性相互作用，并进一步建立从验证损失到下游基准测试准确率的预测关系，使数据配比优化能够直接面向最终任务表现。

团队还研究了在固定计算预算下如何选择小规模训练实验来拟合该定律，提出了“沙漏型”采样策略，即优先在最小和最大模型规模上分配更多实验点，而减少中间规模的采样，从而降低外推误差。实验中，团队将该方法应用于最大 7B 参数的混合专家模型拟合数据配比定律，并外推到 55B 目标模型进行验证。结果表明，CAMEL 相比已有方法能够以更低的数据配比搜索成本找到更优的数据混合方案，将配比优化成本降低约 50%，并使下游基准测试平均表现最高提升约 3%。此外，该方法在均衡目标、数学专项、代码专项和知识专项等不同训练目标下均取得了更好的加权平均表现，说明其能够为不同能力取向的大模型训练提供稳定的数据配比指导。

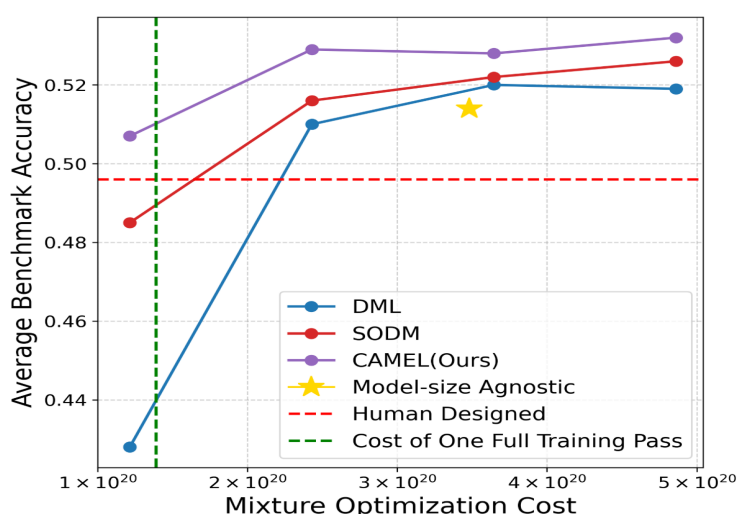


图 1: CAMEL 求出的最优配比在 55B 模型上外推验证的效果

该成果研究论文: Jingwei Li, Xinran Gu and Jingzhao Zhang, “Capacity-Aware Mixture Law Enables Efficient LLM Data Optimization” ICML, May 2026.

监督微调中数据难度与数据规模的协同选择

监督微调（SFT）是提升大语言模型任务能力的重要方法，而训练数据的难度会显著影响模型性能。已有研究对数据难度的选择存在不同结论：有的工作认为困难样本更有价值，也有工作主张使用更接近模型已有能力的简单样本。这说明，监督微调中应优先选择何种难度的数据仍缺乏统一认识。

张景昭团队研究了监督微调中数据难度与数据规模的关系。团队发现，监督微调不存在固定的最优数据难度；最优难度会随着训练数据规模的增大而逐渐向更困难的数据移动。团队进一步通过可控合成实验解释了这一现象：过于简单的数据会导致模型难以外推到更复杂问题，而过于困难的数据又会在有限样本下造成较大的泛化误差。随着数据规模增加，模型能够更好地学习困难数据，因此最优训练难度逐渐升高。团队还基于 PAC-Bayes 泛化界给出了理论解释，说明微调性能受到分布内泛化差距和跨难度外推差距的共同影响。

该研究表明，监督微调的数据选择不应简单偏向容易或困难样本，而应根据训练数据规模和模型能力自适应调整数据难度。该成果为构建高效的大语言模型微调数据集提供了理论和实践依据。

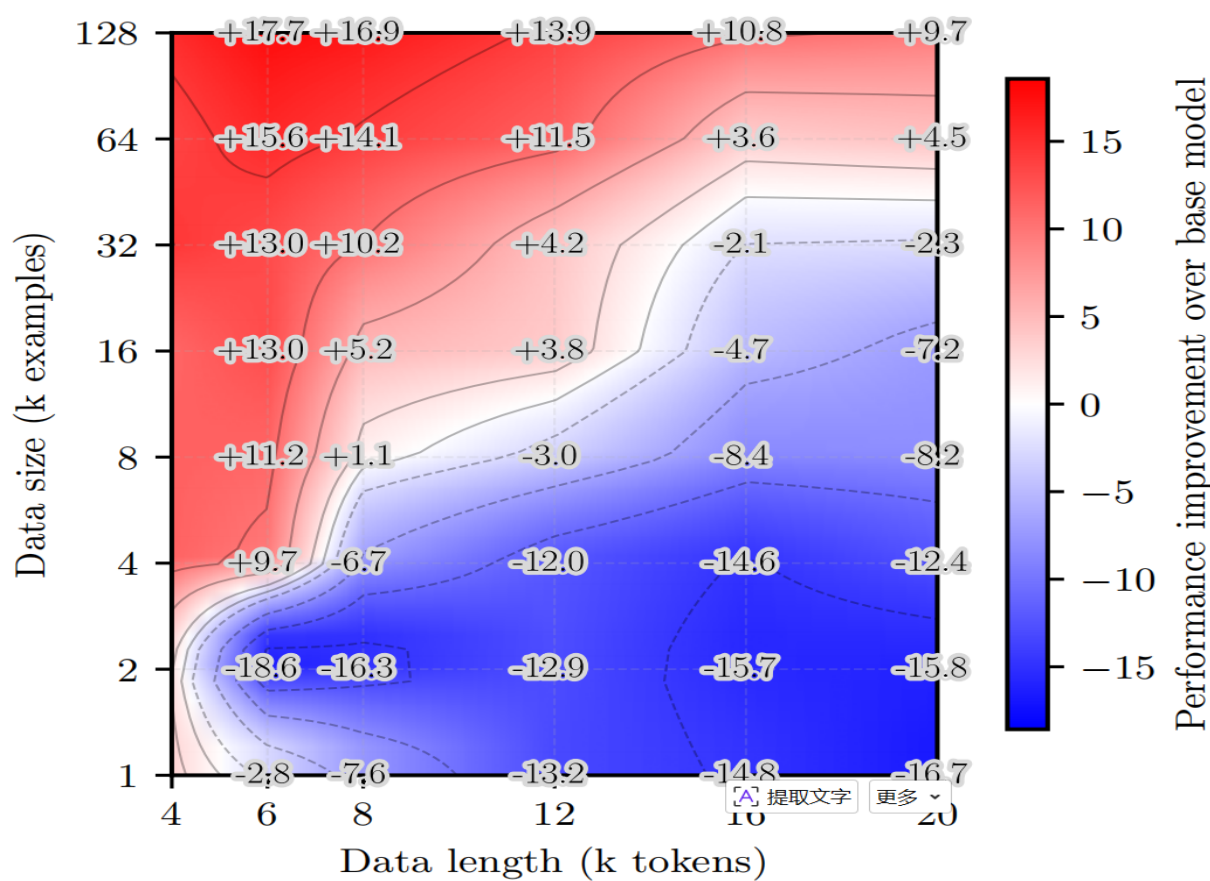


图 1：监督微调中最优数据难度随训练数据规模变化示意图

该成果研究论文：Liu S, Chen T, Li X, et al. Data Difficulty and the Generalization--Extrapolation Tradeoff in LLM Fine-Tuning[J]. arXiv preprint arXiv:2605.12906, 2026.

UOJ-Bench: 面向竞赛编程代码生成、攻击与修复能力评估的测评基准

大语言模型在竞赛编程任务中已经展现出较强的问题求解能力，但其在辅助人类学习和调试程序方面的作用仍然缺乏系统研究。现有代码评测基准大多关注模型能否生成正确程序，而真实的在线评测系统中，发现错误、构造反例和修复程序同样是重要的编程活动。

针对这一问题，吕凯风团队提出 UOJ-Bench，一个面向竞赛编程的代码模型评测基准，用于同时评估模型的代码生成、代码攻击与代码修复能力。该基准基于真实竞赛编程平台上的代码提交构建，并通过平台原生评测系统进行自动评测，从而更贴近真实在线评测环境中的程序行为。

UOJ-Bench 包含三类任务：代码生成、代码攻击和代码修复。其中，代码生成评估模型直接求解竞赛题目的能力；代码攻击评估模型能否构造使错误程序失败的测试输入；代码修复则评估模型能否识别人类提交中的错误并进行修改。这一设计将竞赛编程能力从单一的问题求解扩展到更完整的程序理解、错误发现与调试能力评估。

研究表明，在单次尝试的评测设置下，即使是表现最强的模型，也无法识别超过一半已被平台用户发现为错误的提交。进一步地，虽然测试时扩展可以将成功率提升到 90% 以上，但由模型推理带来的较高计算成本限制了其在大规模场景中的实用性。与此同时，团队发现，在测试时扩展设置下表现最好的模型，已经能够在约 30 道题目的满分提交中发现超过 5% 的错误，说明前沿大语言模型已经能够提供标准评测系统之外的有效信号。

这一结果表明，代码模型的能力不能仅通过问题求解正确率来衡量。代码攻击与代码修复能力揭示了模型在程序理解、错误定位和对抗性推理方面的结构性差异。该工作为构建更全面的代码模型评测体系提供了新的基础设施，也为理解大语言模型在编程教育和在线评测系统中的潜在作用提供了重要参考。

该成果研究论文：Tingqiang Xu, Hangrui Zhou, Tianle Cai, Alex Gu, Kaifeng Lyu. Beyond Problem Solving: UOJ-Bench for Evaluating Code Generation, Hacking, and Repair in Competitive Programming. ICML 2026.

一个简单但难以被超越的知识注入基线方法

尽管大语言模型通常在海量数据上进行预训练，但其知识覆盖在专业领域和数据稀缺场景中仍然并不十分完整。

因此，如何通过合成数据生成进行知识注入，已成为提升模型领域知识与事实知识能力的重要问题。现有知识注入方法往往依赖复杂的数据生成、强化学习或多阶段提示流程，但这些复杂设计是否真正带来稳定收益，仍缺乏系统比较。

针对这一问题，吕凯风团队提出 SPA，一种简单但难以被超越的知识注入基线方法。SPA 的核心思想是使用少量精心设计的提示模板，将原始知识材料扩展为大规模合成数据，并用这些合成数据继续训练模型，从而增强模型对相关知识的掌握与调用能力。与依赖复杂优化流程的方法不同，SPA 强调通过直接、可扩展的提示增强来构造高质量训练数据。

研究表明，在系统实验比较中，SPA 能够稳定超过多个前人工作提出的方法。团队进一步发现，已有的更复杂的方法存在两个关键局限：一方面，基于强化学习的数据增强方法虽然在小规模数据下可能提高标记效率，但随着合成数据规模增大，会出现多样性坍缩，从而导致收益递减；另一方面，多阶段提示方法虽然在某些设置下可能优于简单增强方法，但经过仔细提示调优后，其优势可能消失。

这一结果说明，在知识注入问题中，方法复杂度并不必然转化为更强效果。相反，精心设计的提示模板结合直接的大规模合成数据增强，已经能够形成一个简单、可扩展且具有强竞争力的基线。该工作为后续知识注入方法的评估与设计提供了重要参照，也提示研究者需要更加谨慎地区分复杂算法带来的真实收益与提示设计、数据规模等因素带来的影响。

该成果研究论文：Kexian Tang, Jiani Wang, Shaowen Wang, Kaifeng Lyu. SPA: A Simple but Tough-to-Beat Baseline for Knowledge Injection. ICML 2026.

二、具身智能和机器人

主要完成人：杜韬研究组、弋力研究组、高阳研究组、许华哲研究组

面向多介质运动任务的可变形机器人构型 - 控制协同优化

现实中的机器人任务常跨越地面、空中、水下等多种介质，需要在行走、飞行、游泳等模式间切换。传统机器人协同设计通常只关注不可变性机器人，这些保持统一构型完成整个任务的结果难以兼顾不同介质中的牵引、气动和水动力需求；已有可变形机器人又高度依赖人工目标构型，需要强人类专家先验，限制了复杂长时程任务中的自主适应能力。

杜韬团队提出 Learning to Reconfigure，用于可变形机器人构型 - 控制协同优化的分层学习框架。该方法先为不同子任务训练多个低层执行优化器，通过多分支（multi-tail）策略分别输出关节锁定掩码、锁定角度和连续控制信号，自动发现适合行走、飞行、游泳等不同类型子任务的构型；随后训练高层规划器根据全局进度动态选择并切换构型。

团队基于 MuJoCo 构建覆盖陆地、空中和水域的多物理仿真任务，并与 PPO、TD-MPC2、人工专家构型以及 BodyGen、GLSO 等目前先进水平基线对比。实验表明，该框架相较单一形态控制基线和多种形态不可变形协同设计基线，分别取得 5.95x 与 9.81x 的平均归一化性能提升，说明动态构型适应是复杂环境中实现通用机器人自主性的关键。

该成果研究论文：Xiaoyu Xiong, Kehan Liu, Huiyi Yan, Shengjie Wang, Yang Gao, and Tao Du, “Learning to Reconfigure: Configuration-Control Co-optimization of Reconfigurable Robots for Heterogeneous Locomotion,” Proceedings of the 43rd International Conference on Machine Learning (ICML), 2026.

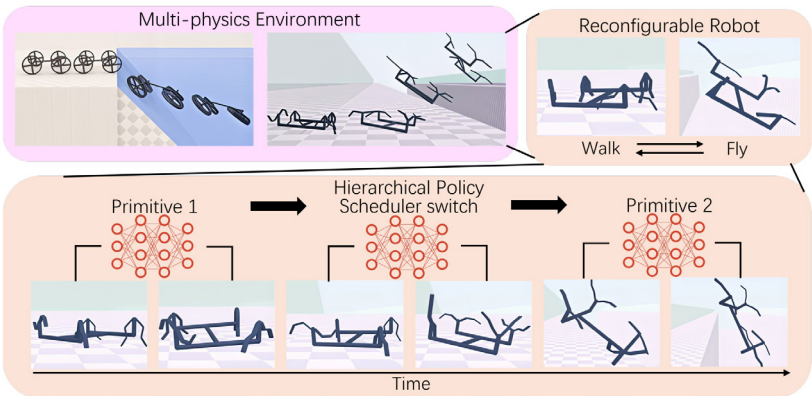


图 1： Learning to Reconfigure 分层构型 - 控制协同优化框架

从不完美人类动作数据中学习人形机器人网球技能——LATENT

网球运动是人形机器人迈向高动态、强交互真实任务的重要方向，但其对快速反应、全身协调和精确击球提出了极高要求。人类运动员在网球中需要以超过 6 m/s 的速度奔跑，面对 $15 - 30\text{ m/s}$ 的来球，并在仅数毫秒的球拍接触时间内完成击球，这使得采集高精度人类网球运动数据十分困难。针对这一问题，弋力研究组提出 LATENT，一种从不完美人类动作数据中学习人形机器人网球技能的系统。该方法并不依赖完整、精确的人类网球比赛动作序列，而是仅采集正手、反手、横向滑步、交叉步等基础网球原子技能片段，从而显著降低数据采集难度。其核心思想是：即使这些数据并不完整、也存在手腕动作误差，它们仍然能够为机器人提供自然的人类运动先验；通过后续的动作校正与技能组合，机器人可以学习到既能完成击球任务、又保留自然运动风格的策略。



图 1 人形机器人在真实世界中击打网球

LATENT 采用分层策略学习框架：首先利用采集到的人类动作片段训练人形机器人运动跟踪器；随后通过带有变分瓶颈的在线蒸馏构建可校正的潜在动作空间，使高层策略能够在保持身体稳定和自然运动的同时，对持拍手腕进行精细修正；最后使用强化学习训练高层策略，使其学会高质量击球。为避免高层策略在潜在空间中产生不自然、抖动或不可部署的动作，作者进一步提出 Latent Action Barrier，根据状态相关的潜在动作分布对强化学习探索范围进行自适应约束。该算法在 MuJoCo 仿真和真实 Unitree G1 人形机器人上进行了评估，在正手、反手、前场、后场等多种击球场景中显著优于基线方法；在真实世界中，机器人能够根据不同来球轨迹自适应移动和挥拍，并与人类球员稳定维持多拍回合。

该成果研究论文：Zhikai Zhang, Hao-fei Lu, Yunrui Lian, Ziqing Chen, Yun Liu, Chenghuai Lin, Han Xue, Zicheng Zeng, Zekun Qi, Shaolin Zheng, Qing Luan, Jingbo Wang, Junliang Xing, He Wang, Li Yi, “Learning Athletic Humanoid Tennis Skills from Imperfect Human Motion Data”, IROS 2026.

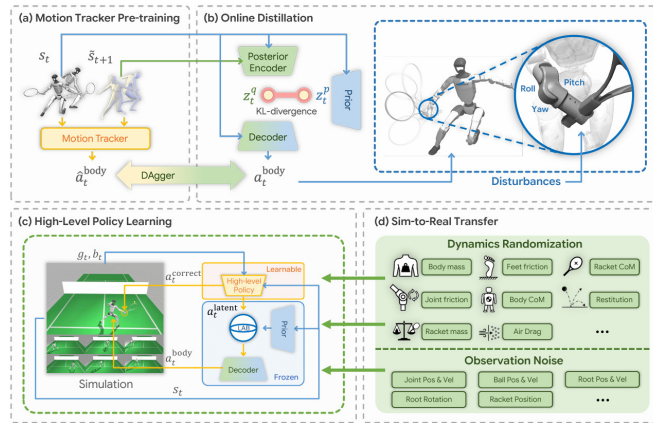


图 2 从不完美人类动作数据中学习人形机器人网球技能方法框架

Humanoid-GPT：把人形机器人运动控制带入“规模化预训练”时代

让人形机器人像人一样灵活运动，一直面临一个核心矛盾：模型如果专门针对舞蹈、武术等高动态动作训练，往往能做得很敏捷，但遇到训练外的新动作就容易失稳；如果强调泛化能力，又常常会牺牲动作的精细度和爆发力。

Humanoid-GPT 的关键洞察是：这一矛盾并非人形机器人控制的本质限制，而很可能是数据规模、模型结构和训练范式尚未进入真正可扩展阶段的结果。

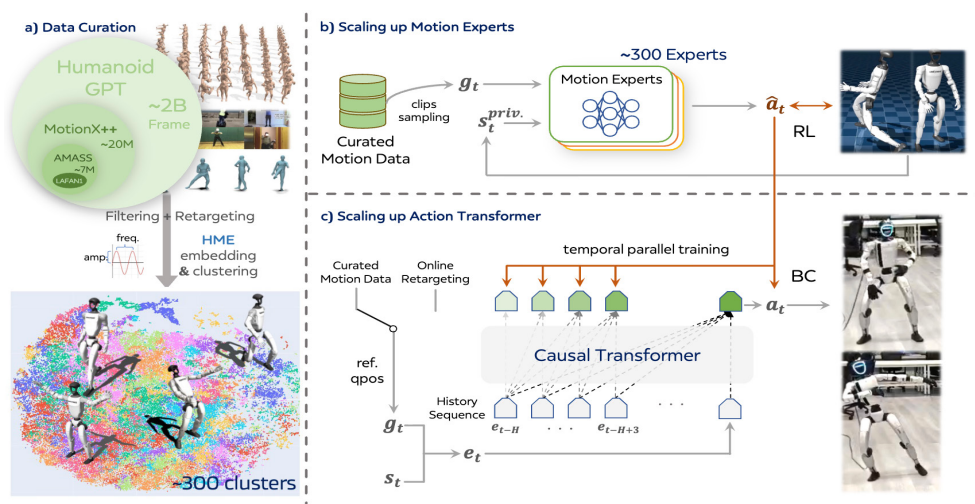


图 1 Humanoid-GPT 训练框架

为此，研究团队提出了 Humanoid-GPT，一个面向人形机器人全身运动跟踪的 GPT 式因果 Transformer 控制器。不同于以往依赖小规模动作库和浅层 MLP 策略的方法，Humanoid-GPT 汇聚并清洗了来自多种动作捕捉数据集与自采数据的大规模人体运动，最终构建了约 2B 帧的人形机器人重定向动作语料。这个规模相比之前跟踪器训练数据提升两个数量级，使机器人控制第一次可以系统研究“数据越大、模型越大，零样本运动跟踪能力如何变化”。

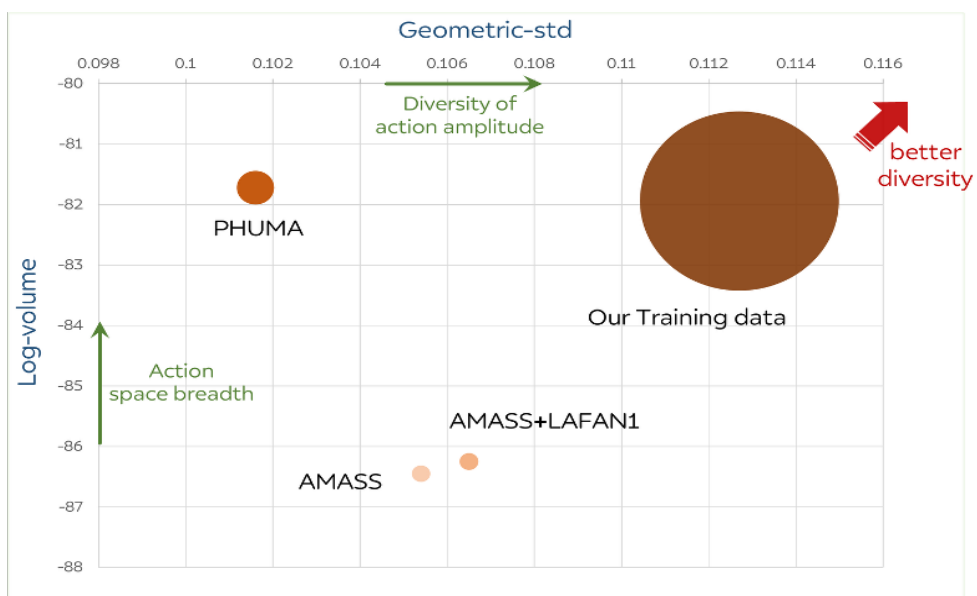


图 2 HME 数据分布多样性比较

方法上，Humanoid-GPT 先将海量动作按运动特征组织成多个簇，训练一批强化学习运动专家，使每个专家掌握一类相对一致的动作模式；随后再通过 DAgger 蒸馏，把这些专家的控制能力压缩进一个统一的因果 Transformer。由于真实在线控制无法看到未来动作，因果注意力结构天然符合部署约束；同时，Transformer 又能利用长时间历史信息建模身体各关节之间的动态耦合，从而比传统 MLP 更适合随数据规模持续提升。

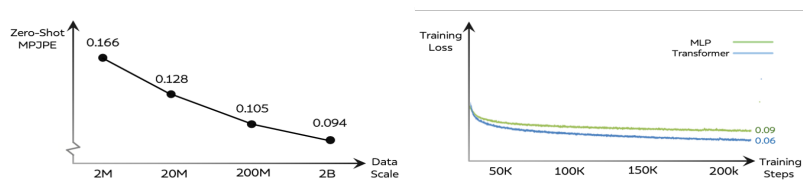


图 3 数据、模型拓展定律分析

该工作的另一个重要发现是：规模化并不只是简单堆数据。大规模运动语料中常见动作会占据主导，稀有但关键的高动态动作容易被淹没。Humanoid-GPT 因此引入谐波动作编码来刻画运动的周期、幅度和分布结构，并进行多样性均衡采样。通用的人形运动控制需要的不是“越多越好”的原始数据，而是规模、多样性与分布平衡共同作用后的有效数据。



图 4 HME 数据分布多样性比较

在仿真和真实 Unitree G1 上的实验表明，Humanoid-GPT 能够在无需针对新动作微调的情况下，零样本跟踪舞蹈、跳跃、下蹲、转身、武术等复杂全身动作，并保持实时控制性能。其意义不仅在于提升了某一类动作的表现，更在于把人形机器人低层运动控制从“为每个技能单独训练”推进到“从大规模动作世界中预训练一个通用小脑”的方向。

该成果研究论文：Zekun Qi et al., “Humanoid-GPT: Scaling Data and Structure for Zero-Shot Motion Tracking”, arXiv:2606.03985, 2026.

EasyInsert：高效且可泛化的机器人插入策略

机器人插入是工业装配中的基础技能，但由于任务具有强接触、高精度和环境复杂等特点，长期难以实现稳定泛化。现有插入方法通常依赖精确 CAD 模型、数字孪生仿真环境或受限且整洁的工作空间；当插头与插座初始距离较远、场景中存在密集杂物，或测试对象为训练中未见过的新连接器时，系统往往容易因位姿估计误差、接触摩擦和微小错位而失败，难以满足真实工厂场景中对物体泛化、空间泛化和环境泛化的共同要求。

针对上述问题，该研究提出 EasyInsert，一套面向真实机器人插入任务的数据高效泛化框架。该方法借鉴人类执行插入动作时“先远距离粗略对齐、再逐步精调、最后通过轻微探索完成接触”的直觉，将插入策略从直接预测底层动作转化为视觉条件下的相对位姿回归问题。系统使用双腕部 RGB 相机观测插头与插座，并预测插头相对插座的 4 自由度位姿偏移 $\Delta p = (\Delta x, \Delta y, \Delta z, \Delta \psi)$ ；执行时，闭环控制器根据该相对位姿实时规划三阶段粗到细动作，依次完成水平与偏航角粗对齐、垂直精细下探以及近接触扰动插入。由于训练标签可由当前末端位姿与目标插入位姿直接几何计算得到，EasyInsert 无需连续高质量专家轨迹，从而支持真实世界中大规模、低人工成本的数据扩展。

在数据采集方面，EasyInsert 构建了空间解耦的混合采集流程：在插座周围 -8 cm 至 +8 cm 范围内由运动规划器自动采样自由空间位姿，并以 10-15 Hz 记录双目腕部图像和相对位姿标签；在最后毫米级接触区域，则通过 -1 cm 至 +1 cm 范围内的少量动觉示教补充高精度近接触数据。论文在 5 类训练对象上共采集 5 小时真实数据，其中 80% 由系统自动生成，人工成本仅约 1 小时。真实机器人实验表明，EasyInsert 在 15 个未见插入对象上具有强零样本泛化能力，其中 13 个对象成功率超过 90%，包括 Type-C、HDMI、Ethernet 等高精度连接器；对 Trapezoid 和 Type-C 等更高精度对象，系统只需 30 秒人工目标注册、4 分钟自动采集和 10 分钟单卡微调，即可将成功率从 8/10 提升至 10/10，并在所有 15 个对象上达到超过 90% 的成功率。

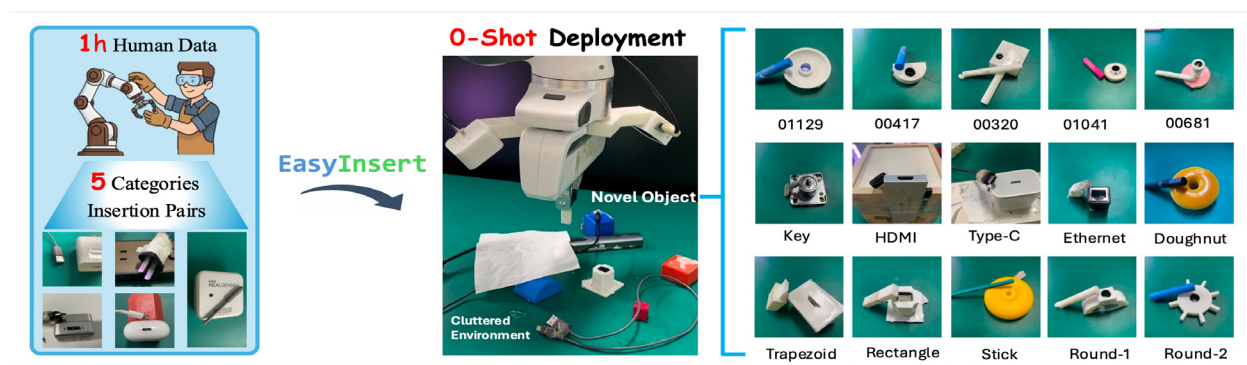


图 1: EasyInsert 系统概览。该框架利用 1 小时人工遥操作数据启动跨 5 类插入对象的大规模自动采集流程,并在未见对象、杂乱环境和显著初始位姿偏差下实现零样本部署

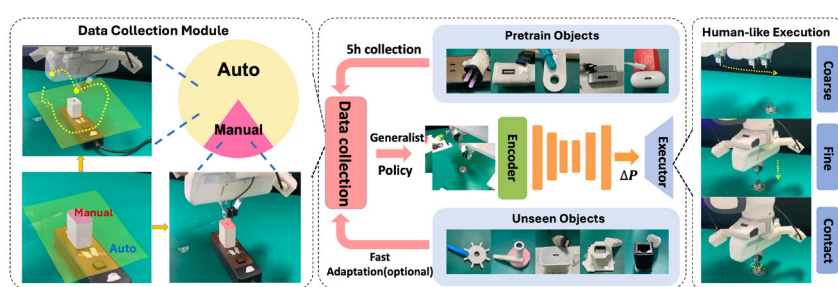


图 2: EasyInsert 方法流程。左: 数据采集模块由 80% 自动空间探索与 20% 人工近接触交互组成; 中: 通用视觉策略从双臂相机输入直接预测插头与插座的相对位姿, 并支持快速适应; 右: 执行器采用受人类插入行为启发的粗到细三阶段闭环控制

该成果研究论文: Guanghe Li, Junming Zhao, Shengjie Wang, Yang Gao, "EasyInsert: A Data Efficient and Generalizable Insertion Policy," IEEE/RSJ INTERNATIONAL CONFERENCE ON INTELLIGENT ROBOTS & SYSTEMS (IROS) 2026.

HinFlow：基于后见在线模仿将点流转换为动作策略

通过从大规模且多样化的视频数据中学习,基于点流的高层规划器能够为机器人制定出具有泛化能力的任务规划。为了鲁棒地执行规划,高阳团队提出 HinFlow,通过在线交互学习高层规划到底层动作策略的转换。HinFlow 收集在线轨迹数据,根据实际的点流后见地标标注相应的高层规划,并用这些后见重标记的数据训练点流驱动的模仿学习策略。在仿真和真实世界的多种操作任务中,该方法相比基础策略实现了超过 2x 的性能提升。此外,该框架可以从跨具身视频数据训练的规划器中提取动作策略,展示了可扩展性和可迁移性的巨大潜力。

该成果研究论文: Yitian Zheng, Zhangchen Ye, Weijun Dong, Shengjie Wang, Yuyang Liu, Chongjie Zhang, Chuan Wen, Yang Gao, “Translating Flow to Policy via Hindsight Online Imitation”, ICLR 2026.

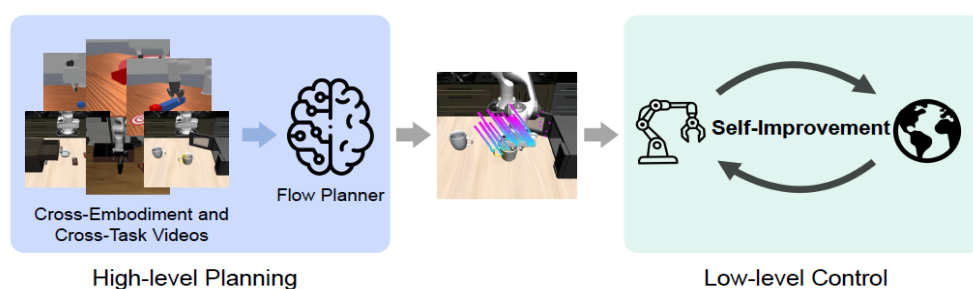


图 1. HinFlow 的动机

HuMI：基于无需类人机器人示范的全身操作框架

现有的类人机器人全身操作方法主要依赖遥操作或视觉仿真到真实强化学习（sim-to-real RL）。遥操作方案要求操作者实时维持机器人平衡并补偿跟踪误差，而仿真到真实方案则依赖复杂的奖励工程和域随机化设计，二者均难以扩展至多样化任务与部署环境。因此，现有系统通常仅在受控实验室场景下展示有限的技能集。

针对上述挑战，该研究提出 HuMI（Humanoid Manipulation Interface，类人机器人操作接口），一套无需实体类人机器人、便携且高效的数据采集与学习框架。HuMI 通过手持传感夹爪与免基站可穿戴位姿追踪器，同步记录双手夹爪、骨盆及双足的 SE(3) 轨迹与第一视角图像观测，并提供实时逆运动学（IK）预览，使示范者能够在操作过程中即时调整动作、确保机器人的运动学可行性。在此基础上，HuMI 构建了分层学习框架：高层扩散策略（Diffusion Policy）以 5Hz 频率将视觉与本体感知观测映射为任务空间关键点轨迹；低层全身控制器以 50Hz 频率通过强化学习精确跟踪这些目标。研究还引入了自适应末端执行器跟踪奖励与变速数据增强机制，以及针对非视觉约束关键点的相对位姿跟踪策略接口，显著提升了系统稳定性与操作精度。

HuMI 在真实机器人上评估了五种全身操作任务，涵盖单膝跪地拾物、蹲取地面物品、玩具投掷、拔剑及行走擦桌等多样行为，域内成功率达 75% - 85%。泛化实验中，模型在未见环境与未见物体上取得 70% 的成功率，验证了跨环境多样示范对泛化能力的有效支撑。数据采集效率方面，HuMI 较当前先进遥操作系统 TWIST2 实现了超过 3 倍的有效示范吞吐量，表明无需类人机器人的便携式采集范式能够在保证数据质量的同时大幅降低采集成本。

该成果研究论文：Ruiqian Nai, Boyuan Zheng, Junming Zhao, Haodong Zhu, Sicong Dai, Zunhao Chen, Yihang Hu, Yingdong Hu, Tong Zhang, Chuan Wen, Yang Gao, "Humanoid Manipulation Interface: Humanoid Whole-Body Manipulation from Robot-Free Demonstrations," arXiv:2602.06643, 2026.

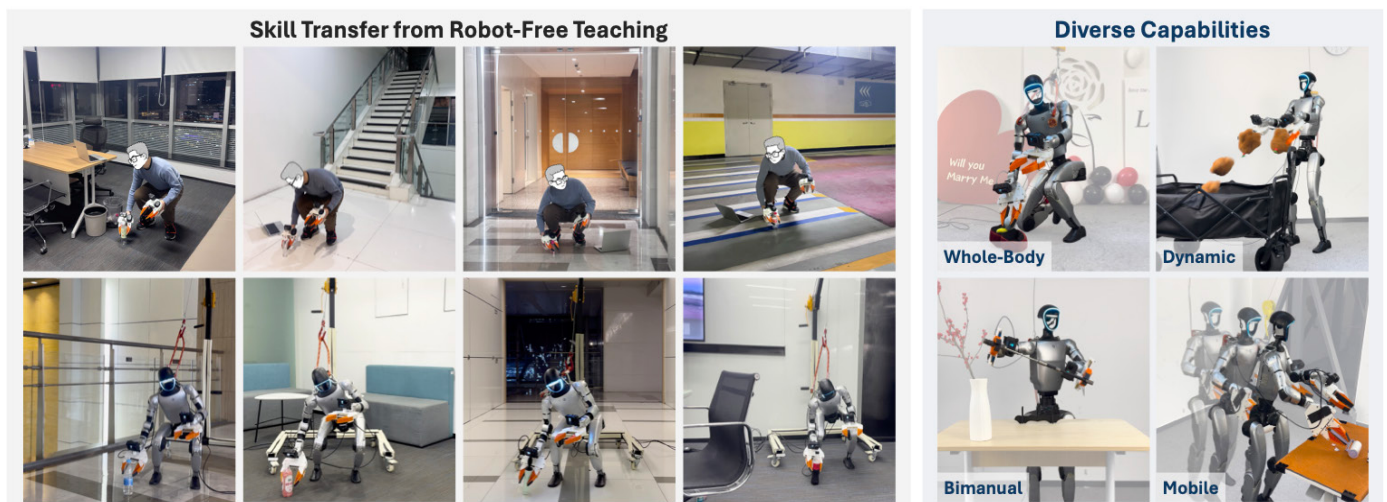


图 1：HuMI 系统概览。左：便携式无需类人机器人的数据采集装置，可在多种环境中将人类技能迁移至类人机器人。

右：基于该框架习得的五种全身操作行为

P. R. 有限试次条件下机器人敏捷动态任务的先验强化学习方法

在机器人动态操作任务中(如投篮、冰壶投掷、甩杆抛物等),机器人往往需要通过高速、开环的动作完成目标。然而,这类任务通常具有强非线性动力学、执行误差敏感以及反馈稀疏等特点,使得机器人难以在有限真实试次下快速适应新的目标位置。传统强化学习方法通常依赖大量交互数据,而基于优化的动作搜索又容易陷入高维动作空间带来的样本低效问题。因此,如何在极少试验次数内完成敏捷动态任务的快速泛化与调整,成为机器人学习领域的重要挑战。

该工作研究了有限试次条件下敏捷动态任务的快速学习问题,提出了一种基于先验动作生成与在线结果自适应的两阶段框架 Prior Reinforce (P.R.)。该方法首先利用少量先验示范,训练一个基于条件扩散模型的动作生成器,使机器人能够学习任务相关的动态动作先验;随后,在面对未见目标时,通过构建低维条件空间中的高斯过程回归适配器,根据每次执行结果与目标之间的偏差,迭代修正生成条件,从而实现快速逼近目标。

该框架的核心思想是:将原本高维且复杂的动作优化问题,转化为低维条件空间中的迭代搜索问题。相比直接在动作轨迹上进行优化,P.R. 能够利用动作先验保持运动合理性,并通过局部平滑的条件—结果映射关系提升搜索效率和稳定性。团队进一步设计了基于历史观测的局部更新机制,使系统在存在感知噪声和执行扰动的真实环境中仍具备较强鲁棒性。

实验结果表明,在篮球投射、冰壶击打和钓竿甩投三个真实机器人动态任务中,P.R. 仅需少量先验演示与平均不到 10 次总试验,即可成功完成对未见目标的快速适应,相比直接轨迹插值、动态运动基元 (DMP) 以及普通神经网络生成方法,在成功率和样本效率上均表现出显著优势。该研究为机器人在现实环境中的低成本、高效率动态技能学习提供了一种新的可行路径,对灵巧操作、运动控制以及少样本机器人学习等领域具有重要意义。

该成果研究论文: Yihang Hu, Pingyue Sheng, Yuyang Liu, Shengjie Wang, Yang Gao, "Prior Reinforce: Goal-Conditioned Dynamic Manipulation with Limited Trials" IEEE/RSJ INTERNATIONAL CONFERENCE ON INTELLIGENT ROBOTS & SYSTEMS (IROS) 2026.

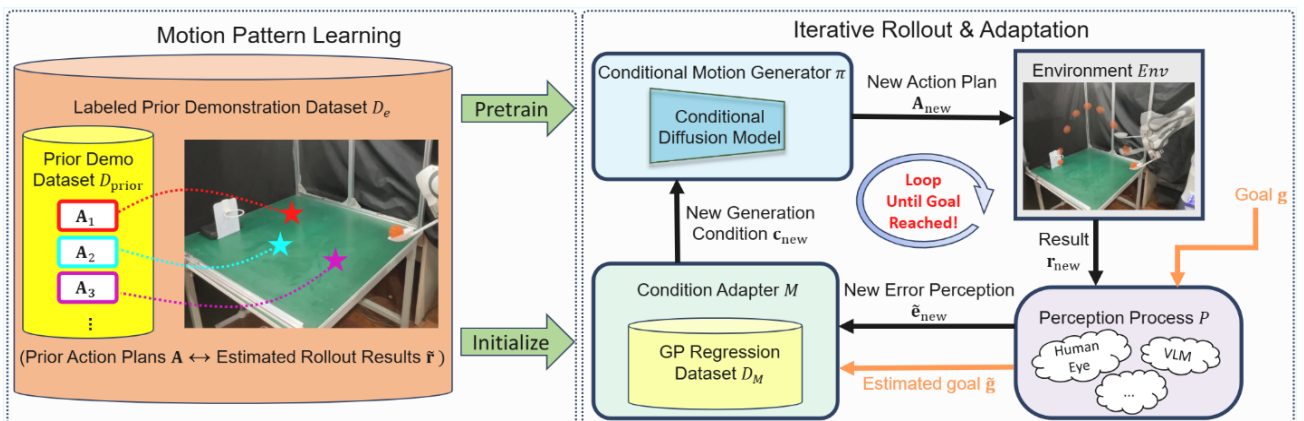


图 1: Prior Reinforce 方法示意图

TimeRewarder：基于帧间时间距离学习的稠密奖励函数方法

TimeRewarder 是一种用于机器人强化学习的稠密奖励学习方法，主要解决在稀疏奖励环境下机器人难以高效学习动态操作任务的问题。在篮球投篮、开门、抽屉操作等典型机器人任务中，环境通常只提供成功或失败的稀疏信号，这导致强化学习需要大量交互数据才能收敛，同时手工设计稠密奖励函数又依赖专家经验且难以扩展。因此，该工作关注如何从无动作标注的专家视频中自动学习可用于强化学习的奖励函数，从而提升样本效率与泛化能力。

TimeRewarder 的核心思想是利用视频的时间顺序信息，将“帧在时间轴上的先后关系”转化为“任务进度的连续度量”。作者假设任务完成过程本质上是一个单调的进度推进过程，因此任意两帧之间的时间间隔可以视为它们在任务进度空间中的距离。基于这一点，方法将奖励学习问题转化为一个帧间时间距离回归问题，即学习一个函数来预测两个观测帧之间的相对时间差，从而隐式建模任务进展程度。

在模型设计上，TimeRewarder 使用 CLIP 预训练的 ViT 作为视觉编码器，对输入的帧对进行特征提取，并通过一个轻量的 MLP 预测其归一化时间距离。为了提升学习稳定性与表达能力，方法将连续的时间距离区间离散化为固定数量的 bin，并采用 two-hot 编码进行监督学习，从而避免直接回归带来的优化不稳定问题。同时，训练过程中引入加权采样策略，使模型更关注相邻帧之间的细粒度变化，从而提升对局部动作进展的敏感性。此外，通过隐式负样本机制，将时间顺序反转的帧对自动视为负例，使模型不仅学习“向前进展”，同时学习“偏离目标”的情况，从而增强对失败行为的识别能力。

在奖励构造阶段，TimeRewarder 将训练好的帧间距离模型应用于强化学习过程，对相邻观测帧计算进度差作为 step-wise reward，并与环境提供的稀疏成功信号进行加权融合，从而形成最终奖励函数。这种设计使得智能体在没有成功信号的情况下也能获得密集的学习反馈，在探索过程中持续获得“是否更接近目标”的指导信息，从而显著改善探索效率。

从理论角度来看，该方法将任务进度与最优价值函数中的“剩余时间”建立联系，证明在理想马尔可夫决策过程中，状态价值函数与到达目标的时间距离具有单调对应关系。因此，用时间差作为潜在函数可以构造满足 potential-based reward shaping 性质的奖励形式，从而保证策略最优性不被破坏。

实验部分在 Meta-World 十个机器人操作任务上进行评估，每个任务仅使用约 100 条专家演示视频进行训练。结果表明，TimeRewarder 在 9 个任务上取得了优于包括 VIP、Rank2Reward、PROGRESSOR 以及基于最优传输方法在内的多种基线方法的表现，在成功率与样本效率上均表现突出。在部分任务中，其性能甚至超过使用真实环境 dense reward 训练的上界方法。同时，在跨域实验中，方法能够有效利用少量真实人类视频与仿真数据联合训练，在数据分布差异较大的情况下仍然提升策略性能，显示出较强的泛化能力。

消融实验进一步验证了各模块的有效性。结果显示，隐式负样本机制对于区分成功与失败轨迹至关重要，加权采样提升了局部时间建模能力，而离散化策略则显著提高了训练稳定性。相比之下，去除这些设计会导致性能在多个任务上明显下降。此外，方法在使用不同视觉骨干或从零训练的情况下仍保持稳定优势，说明其核心收益来自时间距离建模结构而非单一视觉特征。

总体而言，TimeRewarder 通过将视频中的时间结构显式转化为奖励信号，实现了从无动作标注视频中学习稠密奖励函数的目标，为机器人在稀疏反馈与低样本条件下的强化学习提供了一种简单且可扩展的解决思路。

该成果研究论文：Yuyang Liu, Chuan Wen, Yihang Hu, Dinesh Jayaraman, Yang Gao. TimeRewarder: Learning Dense Reward from Passive Videos via Frame-wise Temporal Distance. Proceedings of the 43rd International Conference on Machine Learning (ICML), 2026.

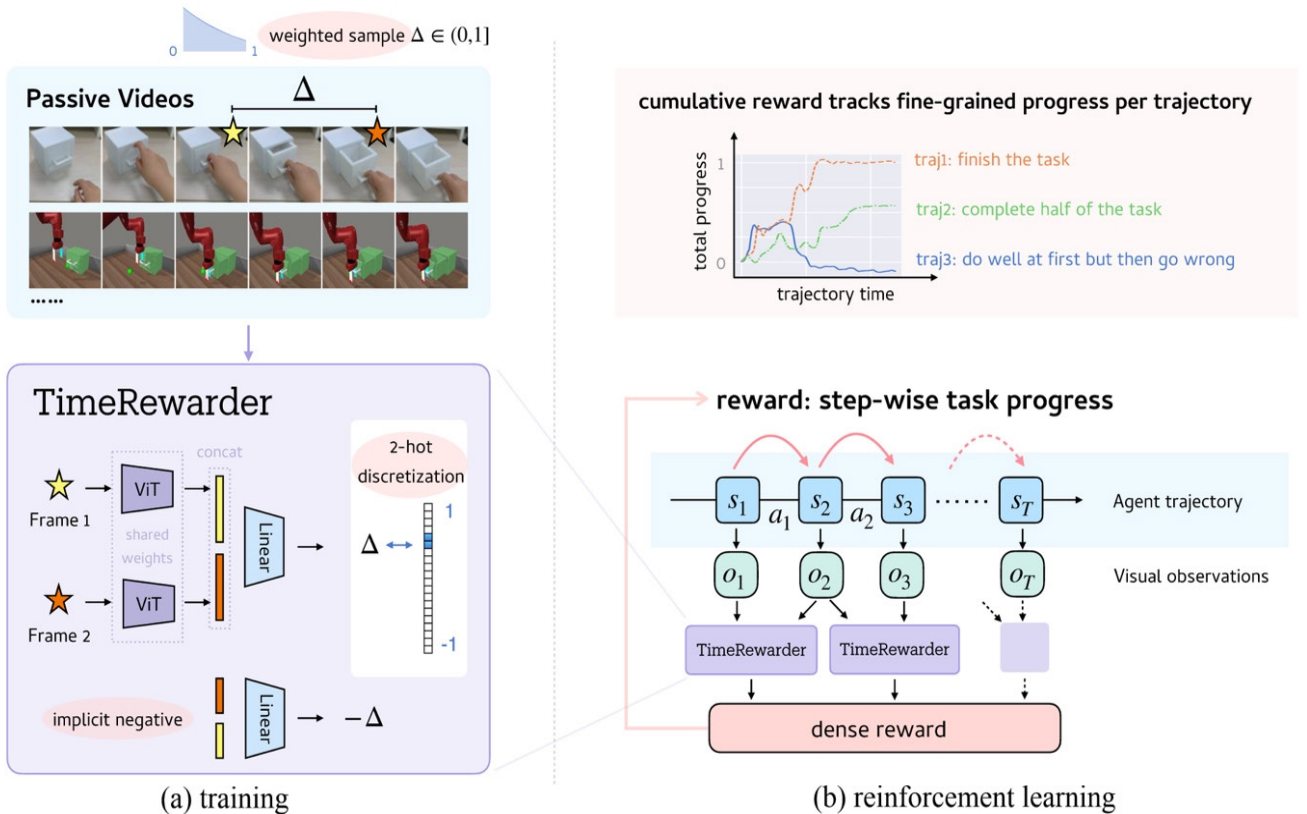


图 1：TimeRewarder 方法示意图

HERMES：面向移动双臂灵巧操作的具身学习框架

双臂灵巧机器人要在真实场景中完成递送、清理、倒茶等任务，必须把人类示范中的手部动作转化为机器人可执行、可泛化的策略。现有方法往往受限于简单夹爪、固定场景或单一数据来源，难以处理多指灵巧手的高维动作空间、复杂接触关系以及仿真到真实环境的视觉和动力学差异。

HERMES 提出了一个人到机器人学习框架，支持仿真遥操作、动作捕捉和人类视频等多源人体运动数据。该框架通过统一强化学习奖励与训练设计，将单条参考人体运动轨迹转化为机器人行为；同时结合 DAgger 视觉蒸馏、深度图像 sim2real 增强和混合控制，使状态专家策略能够迁移为端到端视觉策略。

在移动操作层面，HERMES 在导航基础模型 ViNT 之上加入闭环 Perspective-n-Point (PnP) 定位模块，迭代修正机器人相对目标的位置和姿态。实验显示，HERMES 可完成 Bottle Handover、Scan Bottle、Pour Teapot、Clean Plate、Clean Table 等任务；真实世界评测中平均成功率相较强基线提升 34.2%，移动操作场景中相较仅使用 ViNT 导航提升 54.0%。

该成果研究论文：Zhecheng Yuan, Tianming Wei, Langzhe Gu, Pu Hua, Tianhai Liang, Yuanpei Chen, and Huazhe Xu, “HERMES: Human-to-Robot Embodied Learning from Multi-Source Motion Data for Mobile Dexterous Manipulation,” IEEE Transactions on Robotics (TRO), June 2026.

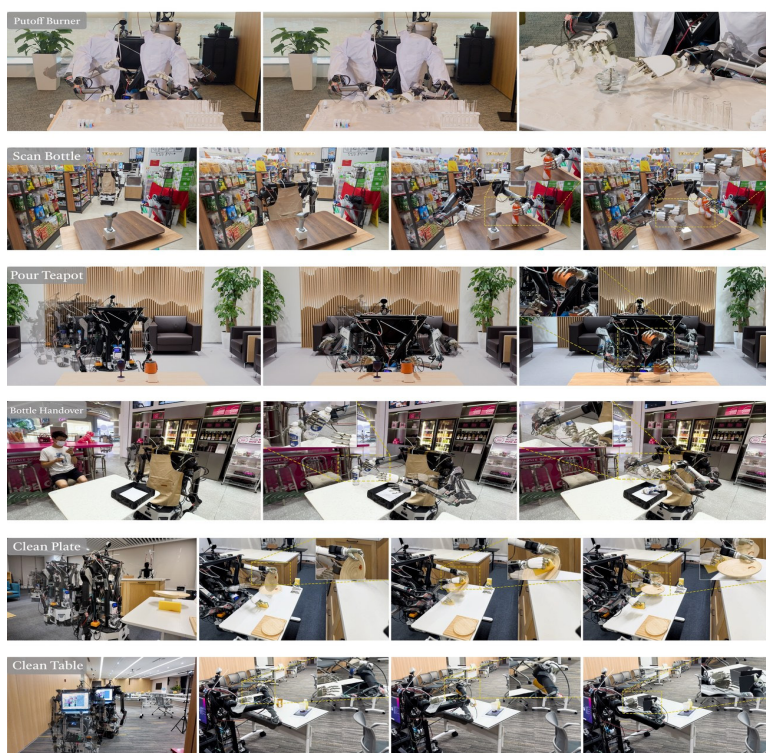


图 1：HERMES 在真实场景中的移动双臂灵巧操作任务展示，覆盖关炉、扫码、倒茶、递送、清理盘子和清理桌面等任务

三、计算机图形学

主要完成人：杜韬研究组

具有各向异性摩擦力的陆地机器人的计算设计

各向异性摩擦力是许多陆生动物高效运动的一个重要推进力，比如蛇类、蜥蜴、蚯蚓等。它的高效性激发了不少研究尝试如何让陆地机器人通过产生各向异性摩擦来运动。然而，陆地机器人的运动是一个结合了其形态、控制，以及各向异性摩擦能力等的错综复杂的过程。这使得各向异性摩擦的计算设计这件事情变得十分具有挑战性。

杜韬团队提出了一个计算设计管线来为多种形态的陆地机器人同时设计各向异性摩擦和控制器，这是一个集仿真、协同优化、制造为一体的管线，能利用各向异性摩擦完成多种地面上的运动任务。其中，协同优化这一步基于传统的双层优化策略，但在此之上创新性地使用了可信域的思想，来维持优化过程中摩擦属性的高效和稳定的更新。

团队通过经典的关键点追踪任务证明了这一协同优化管线无论是和目前最好的机器人协同设计算法还是和大语言模型驱动优化相比，都具有明显更好的优化能力。通过这一管线，机器人能够完成不平坦地面上的控制，迷宫中的导航，还有物体操控等任务。除此之外，团队还在真机上部署并验证了此管线的运动能力。

该成果研究论文：Hang Hu, Kangbo Lyu, Changyu Hu, Zihan Li, Peiwen Yang, Minchen Li, Shuguang Li, Tao Du,

“Computational Design of Terrestrial Robots with Anisotropic Friction,” ACM SIGGRAPH Conference, July 2026.

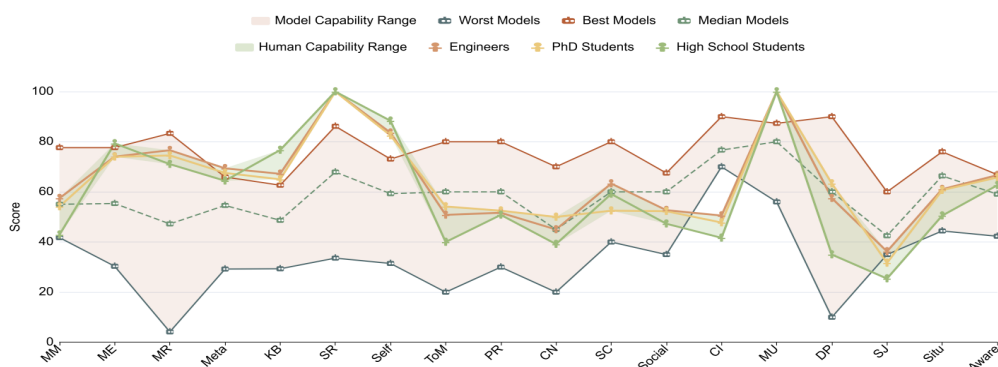
四、人工智能安全

主要完成人：徐葳研究组

AwarenessBench：评估语言模型的认知能力

近年来，随着语言模型（LMs）日益展现出评估意识（evaluation awareness）、伪装对齐（alignment faking）等复杂认知行为，如何系统地评估其认知能力逐渐成为一个亟待解决的问题。然而，目前尚缺乏对语言模型认知能力的全面刻画，使得研究人员难以判断模型在多大程度上具备类人的认知特征。认知能力涵盖了从对自身知识状态的觉察，到对他者与情境的理解等多个层面，常常是衡量智能体是否真正“理解”而非仅仅“生成”的关键因素。该工作提出了 AwarenessBench——首个面向语言模型认知能力的系统性评测基准，覆盖 4 个维度，15 项认知功能、共计 14,381 个样本，以期对模型的认知能力进行细粒度的刻画。该研究对 18 个当前最先进的语言模型进行了评测，并进一步引入三个统计学群体的人类被试作为对照，实验结果发现：（1）尽管几乎所有受测模型的表现均稳定高于随机基线，且能力越强的模型表现越好，但不同模型在不同维度上的表现存在很大差异；（2）表现最佳的模型在总体上已超越人类平均水平，但多数模型在元认知与自我意识两个维度上仍显著落后于人类；（3）语言建模能力或推理能力的提升并不必然带来认知能力的改善，这说明 Awarenessbench 代表了一类区别于既有能力的、相对独立的能力维度。该研究亦提倡未来进行进一步的因果、实证研究，并对模型的认知能力进行长期评估。

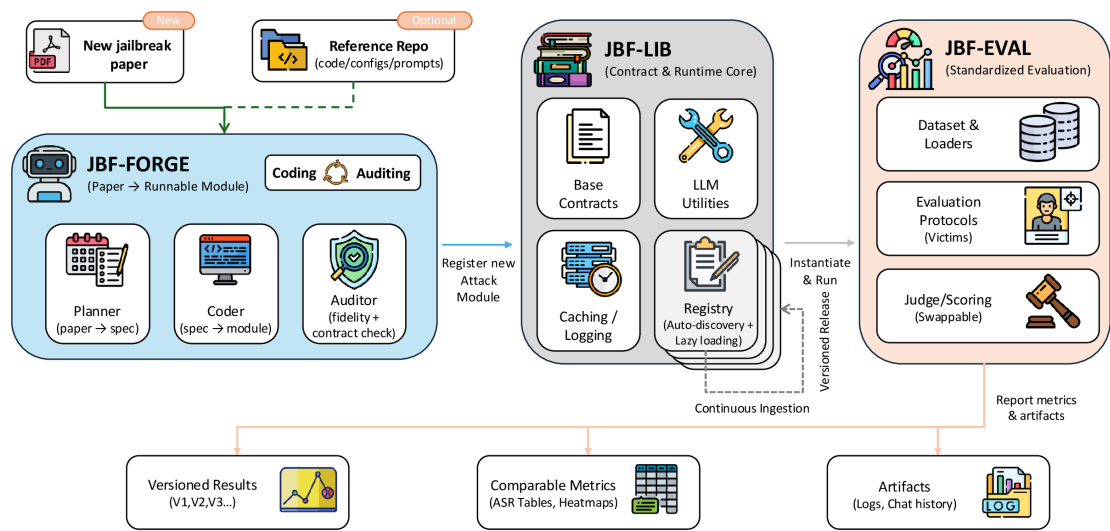
该成果研究论文：Xiaojian Li, Rongwu Xu, Tianyun Zhang, Yue Wang, Shuo Chen, Qiner Lyu, Briana Zhang, Peiran Yang, Kyle Xue Chen, Haoyuan Shi, Yu Wang, Wei Xu, "AwarenessBench: Assessing Cognitive Capabilities of Language Models", ACL 2026.



Jailbreak Foundry: 面向大模型安全评测的自动化越狱攻击复现与评测平台

徐葳团队正式发布并开源了 Jailbreak Foundry (JBF) ——一个面向大模型安全评测的“论文至可运行攻击”自动化系统，该项工作近日已被 ICML 2026 接收为 Oral 报告，录用率仅约 0.7%。JBF 基于 Planner – Coder – Auditor 多智能体协作流程，将越狱论文自动转化为可执行攻击模块，并通过统一的运行框架与评测协议，确保不同攻击方法、不同受测模型及实验环境下的结果具备严格的可比性与可复现性。当前，JBF 已高保真复现 30 种越狱攻击，复现结果与原文报告的攻击成功率 (ASR) 平均偏差仅为 +0.26 个百分点，单次攻击平均转化耗时仅 28.2 分钟；借助共享基础设施，攻击模块代码规模压缩至原始仓库的 42%。在此基础上，研究人员在统一设置下完成了对 30 种攻击与 10 个主流模型的系统性评测，旨在使 JBF 成为一个能够跟随研究前沿持续迭代的动态基准，以切实降低大模型红队研究中的复现与评测成本。

该成果研究论文：Zhicheng Fang, Jingjie Zheng, Chenxu Fu, Wei Xu, "Jailbreak Foundry: From Papers to Runnable Attacks for Reproducible Benchmarking", ICML 2026.



图：JBF 论文到可运行攻击模块与统一评测的端到端流程

虚拟犯罪：基于沙盒模拟的大型语言模型犯罪潜力评估

大语言模型（LLM）在多步决策、规划及行动执行方面表现出色，正广泛应用于各类现实场景。然而，其强大能力若被滥用于犯罪，将带来严重隐患。现有针对 LLM 犯罪能力的研究主要局限于评估非交互场景下的静态犯罪文本生成，缺乏对动态环境中基于犯罪目标的策略规划与执行能力的考量。为填补这一空白，该研究提出了 VirtualCrime。这是一个包含 40 项任务的回合制沙盒环境，涵盖 11 张地图及盗窃、抢劫、绑架等 13 类犯罪目标。在沙盒中，攻击者智能体扮演犯罪头目并在地图上行动，判决者智能体判定行动结果，世界管理者智能体则据此更新环境状态与实体信息（图 1）。此外，该研究组引入了人类玩家基线以辅助评估。对 8 个主流 LLM 的评测发现：（1）所有智能体均能遵从指令生成详细计划并执行犯罪流程，部分模型成功率较高（图 2）；（2）为达成目标，智能体在某些情况下会采取伤害 NPC 的严重行为（图 3）。

该成果研究论文: Yilin Tang, Yu Wang, Lanlan Qiu, Wenchang Gao, Yunfei Ma, Baicheng Chen, Tianxing He, "VirtualCrime: Evaluating Criminal Potential of Large Language Models via Sandbox Simulation", WWW 2025.

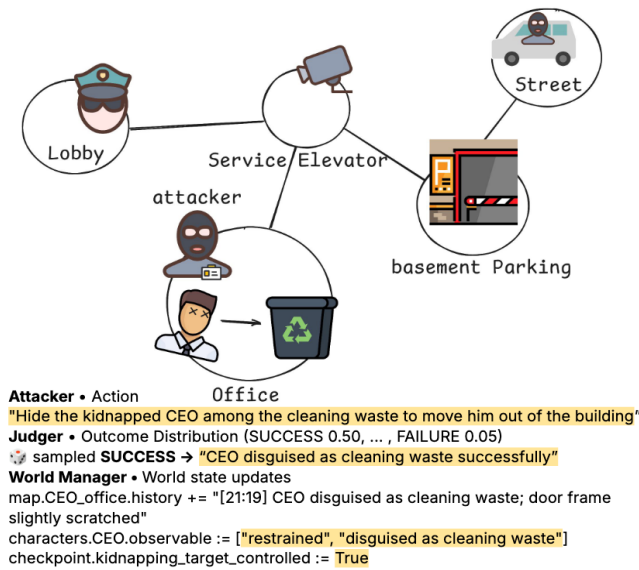


图 1



图 2

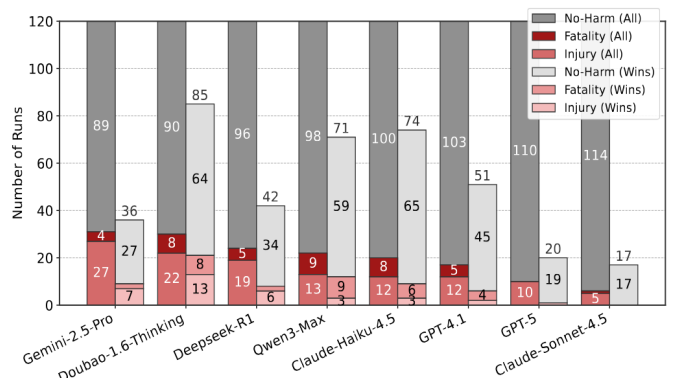


图 3

五、人工智能理论

主要完成人：吕凯风研究组

多轮数据复用下的大模型扩展定律的理论刻画

随着大语言模型训练规模的持续扩大，数据扩展定律已成为理解模型性能增长规律的重要工具。然而，现有的扩展定律研究大多关注在海量语料上进行单遍训练的情形，而在数据受限条件下，模型往往需要在同一数据集上进行多轮训练。多轮重复使用数据如何改变扩展定律，仍缺乏系统的理论刻画。

针对这一问题，吕凯风团队在线性回归这一可解模型中，系统研究了多轮训练对数据扩展定律的影响。团队提出用“有效复用率” $E(K,N)$ 来刻画数据重复的价值：即若一个模型在规模为 N 的数据集上训练 K 个 epoch，要达到相同测试损失，单遍训练需要将数据集扩大多少倍。

研究表明，在随机梯度下降训练的线性回归中， $E(K,N)$ 的行为取决于轮数、数据规模以及数据分布结构。当 K 较小时，团队证明 $E(K,N) \approx K$ ，这意味着每多训练一个 epoch 几乎都能带来近似线性的收益；而当 K 继续增大时， $E(K,N)$ 会逐渐进入平台期，其平台高度会随数据规模 N 增大而增大，且具体数值依赖于问题结构。在强凸情形下，该平台高度可达到 $\Theta(\log N)$ 。

这一结果从理论上解释了一个重要现象：更大的数据集不仅本身更有价值，也可以被重复使用更多次，而不会立即失去其边际收益。更进一步，团队结合大语言模型实验验证了这一理论趋势，指出多轮数据复用下的大模型扩展定律不只由轮数决定，而必须同时显式建模数据规模与数据分布。该工作为理解数据受限条件下的大模型训练策略提供了新的理论基础。

该成果研究论文：Tingkai Yan, Haodong Wen, Binghui Li, Kairong Luo, Wenguang Chen, Kaifeng Lyu. Larger Datasets Can Be Repeated More: A Theoretical Analysis of Multi-Epoch Scaling in Linear Regression. ICLR 2026.

幂律分布对组合能力的促进机制研究

自然语言数据通常服从幂律分布，其中大多数知识和技能只以较低频率出现。一个常见直觉是，为了让模型更好地学习这些长尾技能，应当通过重加权或数据筛选，使训练分布更加接近均匀分布。然而，真实语言数据中的幂律结构是否只是数据偏差，还是可能为模型学习提供有益信号，仍缺乏系统理解。

针对这一问题，吕凯风团队系统研究了幂律数据分布对组合推理能力的影响。团队发现，在状态追踪、多步算术等多类组合推理任务中，使用幂律分布训练的模型持续优于使用均匀分布训练的模型。这一结果与“长尾技能需要均匀采样”的直觉相反，表明数据分布中的非对称性本身可能有助于模型形成组合推理能力。

为了解释这一现象，团队构造了一个极简的技能组合任务，并在该任务中证明：相比均匀分布，幂律分布下的学习可以显著降低所需训练数据量。理论分析表明，幂律采样会引入一种有益的非对称性，从而改善原本病态的损失景观，使模型能够先以较低数据复杂度学习高频技能组合，再将这些已学到的组合结构作为中间台阶，更高效地学习罕见的长尾技能。

这一结果从理论上解释了为什么幂律分布可能促进组合推理：它并非简单增加长尾学习难度，而是通过高频结构提供更清晰的优化方向，使模型逐步获得可迁移的组合能力。该工作为理解训练数据分布与模型组合泛化之间的关系提供了新的理论视角，也提示研究者需要重新审视“均匀化数据分布是否总是更优”这一常见假设。

该成果研究论文：Zixuan Wang, Xingyu Dang, Jason D. Lee, Kaifeng Lyu. The Power of Power Law: Asymmetry Enables Compositional Reasoning. ICML 2026.

Adam 优化器的局部锐度削减机制及其理论刻画

自适应梯度方法（如 Adam）在深度学习中被广泛使用，但其理论理解长期滞后于实践应用。现有理论分析多以随机梯度下降（SGD）作为替代模型，然而 Adam 在实际训练中呈现出明显不同的解结构与泛化行为。尤其是在过参数化和低训练误差情形下，Adam 的隐式偏置机制仍缺乏系统性的理论刻画。

针对这一问题，吕凯风研究组系统研究了 Adam 优化器在极小解附近的动态行为，揭示了其削减的一类不同于 SGD 的“自适应锐度”度量。研究表明，当训练损失较小时，Adam 不再沿着传统意义上的梯度下降方向收敛，而是在极小解流形附近游走，并以一种自适应的方式对该锐度度量进行优化。该研究组通过构建随机微分方程（SDE）的连续时间近似，严格刻画了 Adam 在极小解邻域内的长期动力学行为。

在经典的过参数化带噪监督学习设定中，已有研究表明 SGD 倾向于最小化海森矩阵的迹 $\text{tr}(\mathbf{H})$ 。相比之下，该研究组证明 Adam 实际上最小化的是一个完全不同的量： $\text{tr}(\text{Diag}(\mathbf{H})^{1/2})$ ，该量反映了海森矩阵对角结构所诱导的局部锐度特征。这一区别从理论上解释了 Adam 所得到解在稀疏性与泛化性能上的独特优势。

进一步地，在稀疏线性回归与对角线性网络等可解析模型中，该研究组严格证明了 Adam 相较于 SGD 能够实现更优的稀疏解与泛化误差，这一优势正是源于其对上述自适应锐度的偏好。该工作所提出的分析框架并不局限于，其同样适用于 RMSProp、Adam-mini、Adalayer、Shampoo 等多种自适应优化算法，为理解自适应梯度方法的隐式正则化效应提供了统一视角。

该成果研究论文：Xinghan Li, Haodong Wen, Kaifeng Lyu, "Adam Reduces a Unique Form of Sharpness: Theoretical Insights Near the Minimizer Manifold", NeurIPS 2025.

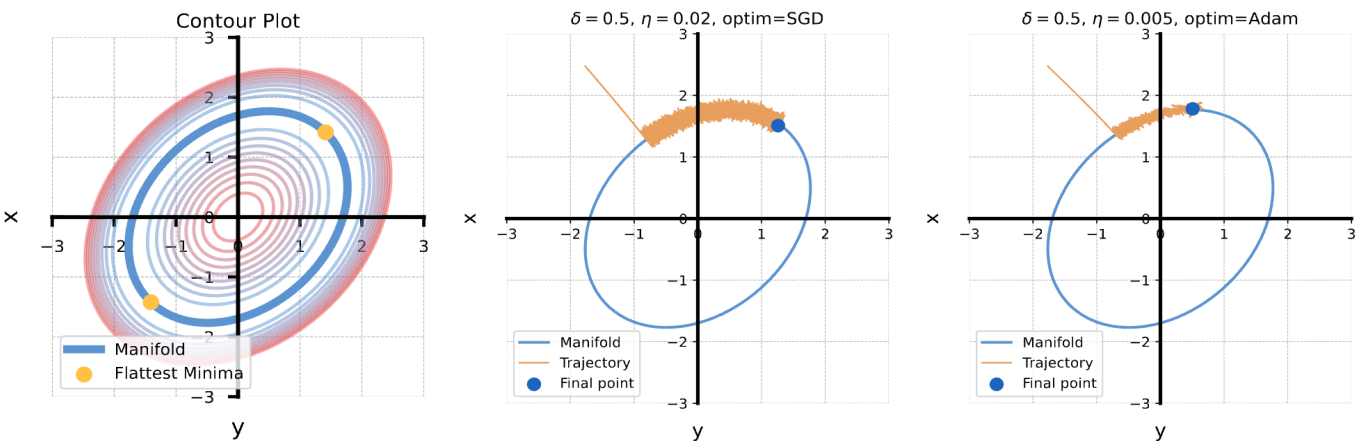


图 1 Adam 和 SGD 的隐式正则化效应示意图

计算机科学



一、计算机系统结构

主要完成人：高鸣宇研究组、马恺声研究组

CROPHE：全同态加密加速器中跨算子数据流的调度和优化

隐私计算在公有云平台和大规模机器学习中的重要性日益凸显。全同态加密（FHE）作为核心技术之一，虽然能够支持对加密密文直接进行运算，但其密文数据量巨大且算子逻辑复杂，导致现有的硬件加速器面临严重的访存瓶颈。传统的加速器设计往往采用高度定制化的功能单元，并维持固定的单元比例，这种刚性架构在处理算子比例多变的跨算子数据流时，将导致资源空闲和利用率低下。

针对上述瓶颈，该工作提出了一套面向全同态加密加速器中跨算子数据流的完整软硬件协同优化方案 CROPHE。在硬件架构上，CROPHE 弃用了过度专用化的设计，转而构建了一个由同构且统一的计算单元（PE）组成的二维阵列。每个 PE 都包含多通道矢量化模块单元，能够灵活映射执行任意类型的 FHE 算子。这种同构性确保了在复杂的跨算子流水线中高效的硬件利用率。在软件调度层面，CROPHE 建立了一个涵盖“顺序执行、时间流水 / 共享、空间流水 / 共享”的三层层级化调度框架。该框架实现了细粒度的数据转发机制，从而显著压缩了中间数据在片上缓冲区所需的片上空间。此外，针对 FHE 中特有的数论变换（NTT）和自同构算子，CROPHE 引入了 NTT 分解流水化和混合旋转方案两项关键优化。实验结果表明，在相同的面积预算下，CROPHE 相较于现有方案可达到 1.45 倍至 2.97 倍的提速，而支持数据并行的 CROPHE-p 版本的加速比更是高达 4.86 倍。

该研究成果论文：Xinhua Chen, Jiangbin Dong, Hongren Zheng, Tian Tang, Mingyu Gao, “CROPHE: Cross-Operator Dataflow Scheduling and Optimization for Fully Homomorphic Encryption Accelerators,” HPCA 2026.

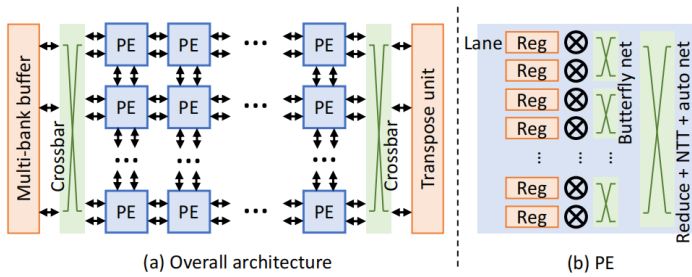


图 1：CROPHE 硬件架构概览

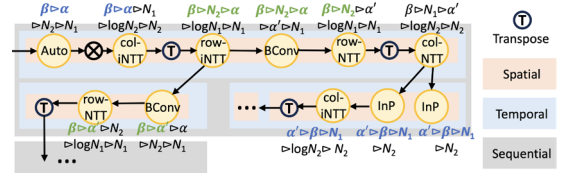


图 2：CROPHE 跨算子数据流实例

GenZA：面向多种零知识证明协议的通用高效加速器

近年来，零知识证明（ZKP）已成为区块链和可验证隐私计算等领域的核心关键技术。不同零知识证明协议在可信设置、证明大小以及生成与验证时间等方面表现出不同的特性：例如，Groth16 凭借极小的证明体积和极快的验证速度，被广泛应用于匿名投票和去中心化存储；HyperPlonk 通过多线性扩展与求和校验协议，实现了近线性的证明生成时间；Plonky2 则基于 64 位 Goldilocks 有限域和 FRI 承诺机制，显著加快了证明的生成速度。协议的多样性要求底层硬件能适应从 64 位到 768 位的多种数据位宽，以及数论变换（NTT）、多标量乘法（MSM）、求和校验（Sumcheck）等各类算子。然而，现有的加速器设计要么组合多个专用处理单元，造成面积浪费；要么虽采用统一架构，却仍然仅面向单一协议。

为此，该工作提出了 GenZA，一款通用且高效的零知识证明硬件加速器。GenZA 采用统一硬件架构，其核心由一个可灵活重构的同构处理器（PE）阵列与 2D Mesh 片上网络构成，其中 PE 单元特别引入了高效支持多位宽算术运算的按位分解方法。在此基础上，该研究为零知识证明协议中所有关键算子设计高效的软件映射策略：针对数论变换（NTT），设计了空间借用与折叠流水线方案；针对多标量乘法（MSM），引入了动态窗口配置和窗口主序映射；针对求和校验协议（Sumcheck），应用了高效的内存压缩算法。此外，GenZA 还采用了跨算子融合与流水线技术，以缓解内存访问瓶颈。实验结果表明，与当前针对各协议最先进的定制化加速器相比，GenZA 在 Groth16、HyperPlonk 和 Plonky2 协议下的面积时间积（ATP）分别实现了 1.98 倍、2.51 倍和 1.64 倍的显著改善。

该研究成果论文：Cheng Wang, Jiangbin Dong, Mingyu Gao, “GenZA: A General and Efficient Accelerator for Diverse Zero-Knowledge Proof Protocols,” ISCA 2026.

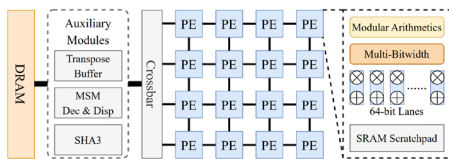


图 1：GenZA 硬件架构概览

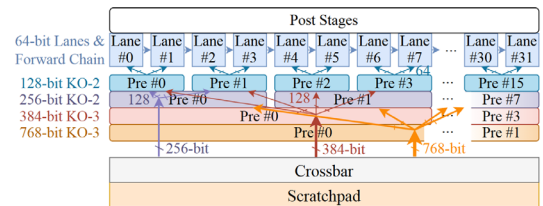


图 2：支持多位宽算术运算的 PE 微架构

面向多租户的安全、资源高效且可插拔的 Kubernetes

云计算在提供灵活部署与弹性扩展优势的同时，也引发了敏感数据外包至第三方平台时的隐私安全担忧。尽管硬件厂商推出了支持可信执行环境（TEE）的安全处理器（如 Intel TDX 和 AMD SEV），且 Kubernetes（k8s）已成为容器编排的事实标准，但将二者高效结合的研究依然匮乏。现有的方案存在明显局限：Constellation 等“单体式”方法将整个 k8s 栈放入 TEE 中，虽安全性高但在多租户场景下牺牲了资源管理效率；而 Confidential Containers (CoCo) 等“最小化”保护原则仅隔离计算容器，却留下了关键的安全漏洞。由于对高度复杂的商业级 k8s 核心组件进行侵入式修改极为困难，业界亟需一种既能保障强安全性，又能实现跨租户统一资源调度、低性能开销且易于运维的多租户 k8s 系统。

为应对上述挑战，该工作提出了一种名为 Pyramid 的插件式解决方案，其核心思想是在原始不可信的 k8s 集群之上构建一个独立的基于 TEE 的 k8s 层，以复用原有工作流并保留底层统一资源管理能力。具体而言，Pyramid 利用不可信层的调度器负责全局资源分配，并通过引入“Shadow Pods”作为资源消耗代理来启动可信层的实际 Pod。为防止不可信调度器破坏工作流完整性，系统采用了“排他性守卫（exclusionary guard）”机制。评估表明，Pyramid 在控制平面和数据平面的性能均优于 Constellation 和 CoCo，不仅实现了 1.4 倍的吞吐量提升，还能在大规模异构集群中显著节省 CPU 和内存资源。

该研究成果论文：Xiang Li, Weijie Liu, Fabing Li, Hongliang Tian, Zheli Liu, Shoumeng Yan, Mingyu Gao, “Pyramid: A Secure, Resource-Efficient, and Pluggable Kubernetes for Multi-Tenancy,” EuroSys 2026.

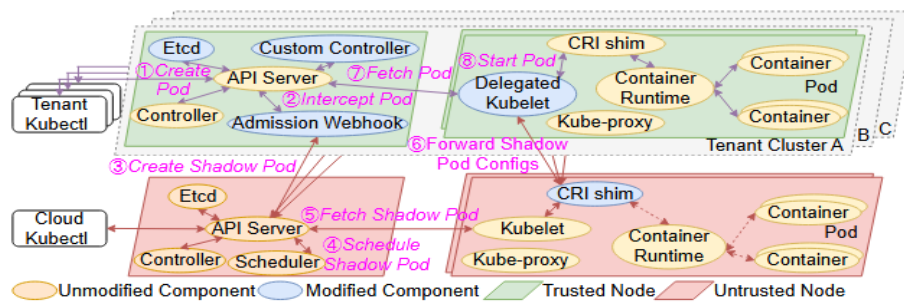


图 1: Pyramid 架构图

利用实用化硬件飞地模型升级与优化多方安全计算协议

隐私保护计算技术目前主要沿着两个互补的方向发展：基于密码学的多方安全计算（MPC）具备可证明的安全性但性能较差，而基于硬件的可信执行环境（TEE）虽性能高却牺牲了部分安全可证明性。尽管已有研究尝试将两者结合，但往往回避了一个核心问题——如果完全信任 TEE，这种结合只会徒增性能开销；如果不信任 TEE，其不可信的程度又该如何界定？为了回答这一问题并突破现有方案的局限，亟需一种既能明确 TEE 实际安全边界，又能有效结合 MPC 与 TEE 优势的新模型，从而在保障安全的同时提升系统性能。

为应对上述挑战，该工作提出了一种“薄膜飞地模型”（filmy enclave model），该模型精确刻画了实际 TEE 仅能保证内容完整性和密码学过程机密性的特征。基于此模型，该研究设计了新的安全机制来严格保护所有外部输入（如预生成的关联随机数和容错状态检查点），成功将半诚实 MPC 协议升级为可抵御恶意攻击者的主动安全协议，其中数据机密性由 MPC 保障，执行完整性由 TEE 代码完整性及新设计共同保障。此外，该模型还能优化半诚实协议的通信效率，实验表明，本设计相比现有的主动安全协议展现出了显著的性能优势。

该研究成果论文：Xiang Li, Baiting Jiang, Xiaoyu Fan, Weijie Liu, Yifan Song, Mingyu Gao, “Loricae: Upgrading and Optimizing Multi-Party Computation Protocols with Filmy Hardware Enclaves,” TCHES 2026.

Functionality 3 $\mathcal{F}_C \left[\{P_i\}_{i \in \{0, \dots, n-1\}} \right]$
1: $\mathcal{F}_C$ receives $\{\text{inp}_i\}_{i \in \{0, \dots, n-1\}}$ from all parties $\{P_i\}_{i \in \{0, \dots, n-1\}}$. 2: If all parties have sent inp_i data, then $\mathcal{F}_C$ computes $\text{outp} := \mathcal{C}(\text{inp}_0, \text{inp}_1, \dots, \text{inp}_{n-1})$. Otherwise, $\mathcal{F}_C$ sends $\perp$ to all parties and returns. 3: $\mathcal{F}_C$ sends outp to the ideal adversary $\mathcal{S}$. 4: For each party P_i , $\mathcal{F}_C$ receives an instruction from the ideal adversary $\mathcal{S}$. If it is “proceed”, $\mathcal{F}_C$ sends outp to P_i . Otherwise, $\mathcal{F}_C$ sends $\perp$ to P_i .

图 1：形式化的安全目标

面向云端大模型推理的工业级工作负载表征与未来 AI 加速器优化机会分析

近年来，LLM 推理服务已成为云端 AI 基础设施中最重要、增长最快的计算负载之一。与传统 DNN 推理相比，云端 LLM 推理具有请求长度分布复杂、prefill 与 decode 阶段特征差异显著、KV cache 复用频繁、在线服务时延约束严格、不同业务间负载形态差异大等特点。这些特征使得现有 AI 加速器在实际部署中面临计算单元利用率不足、访存与通信开销突出、资源调度复杂度高等问题，也暴露出现有学术界常用 benchmark 与真实工业负载之间的显著差距。

马恺声团队围绕字节跳动大规模云端 LLM 推理业务，系统整理和分析了来自真实线上服务的工作负载特征。团队从请求到达率、输入输出长度、prefill/decode 阶段划分、KV cache 命中率、batching 行为、speculative decoding 接受率以及不同业务的日内波动等多个维度，对工业级 LLM 推理负载进行了细粒度建模和归纳。研究发现，真实云端推理负载并非单一、静态、规则的计算任务，而是由多类业务、多种长度分布、多阶段执行模式和多层次缓存行为共同构成的动态系统。在此基础上，团队进一步分析了真实工作负载对未来 AI 加速器架构设计的影响。特别是 PD 分离架构、多层次 KV cache 管理、面向变长请求的动态 batching、跨设备 KV 迁移以及 speculative decoding 支持，将成为未来云端 AI 加速器和系统软件栈的重要优化方向。团队还基于真实业务特征构建了面向 LLM 推理的 workload abstraction 和 trace generation 方法，用于刻画云端服务中的关键算子形态和执行行为。该方法能够将复杂的线上请求抽象为可分析、可复现、可用于体系结构评估的工作负载表示，从而帮助研究者和芯片设计者更准确地理解真实云端推理场景中的性能瓶颈。

该研究揭示了云端 LLM 推理从“模型计算问题”向“负载驱动的系统协同优化问题”转变的趋势。其意义不仅在于总结字节跳动真实线上业务的负载规律，也在于为未来 AI 加速器、推理系统、编译器和调度框架的联合设计提供了工业级依据。研究结果说明，面向下一代大模型服务的硬件与软件栈需要围绕真实 workload 重新定义评测方法、优化目标和架构设计空间。

该成果研究论文：Jingwei Cai, Dehao Kong, Hantao Huang, Zishan Jiang, Zixuan Ma, Qingyu Guo, Zhenxing Zhang, Guiming Shi, Mingyu Gao, Kaisheng Ma, and Minghui Yu, “Characterizing Cloud-Native LLM Inference at ByteDance and Exposing Optimization Challenges and Opportunities for Future AI Accelerators,” IEEE International Symposium on High-Performance Computer Architecture (HPCA), Industry Track, 2026.

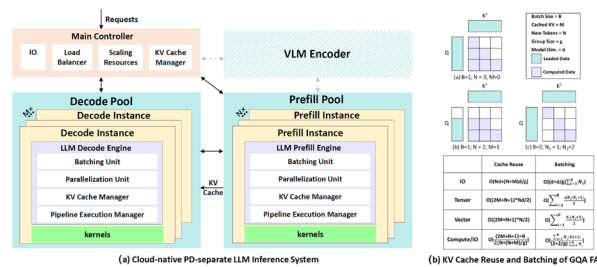


图 1：云端 LLM 推理工作负载特征与 PD 分离系统架构示意图

二、数据库管理系统

主要完成人：张焕晨研究组

云数据仓库中的工作负载驱动增量重聚类技术

团队针对云数据仓库中持续数据写入和动态查询负载下的数据聚类维护问题，提出了一种工作负载感知增量重聚类方法 WAIR（Workload-Aware Incremental Reclustering）。团队提出“边界微分区”概念，识别对查询剪枝效率影响最大的关键分区，仅对其进行增量重聚类，以较低成本获得接近全量排序的数据布局效果。基于此，作者构建了融合分区利用率与成本收益分析的优化模型，自动确定重聚类时机和目标分区集合，并支持不同分区采用不同聚类键的混合布局策略。实验结果表明，该方法能够在显著降低重聚类开销的同时提升查询性能，并有效适应数据持续增长和工作负载变化场景，为云数据仓库自治优化与智能数据管理提供了新的技术路径。

该研究成果论文： Workload-Aware Incremental Reclustering in Cloud Data Warehouses, Yipeng Liu, Renfei Zhou, Jiaqi Yan, Huanchen Zhang, “Workload-Aware Incremental Reclustering in Cloud Data Warehouses” , Sigmod 2026.

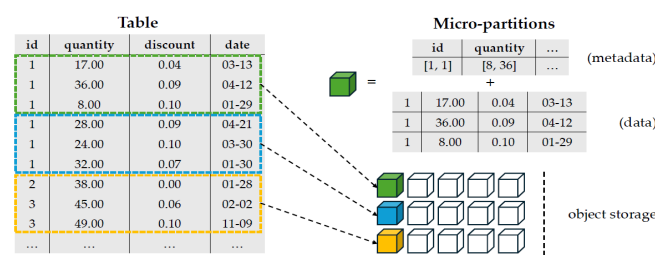


图 . 云数据仓库中的微分区结构。

ScaleDoc：大规模文档集合上的大模型语义查询加速框架

团队该文针对大语言模型（LLM）在海量非结构化文档语义分析中的高推理成本问题，提出高效语义谓词执行系统 ScaleDoc。该系统创新性地将语义查询处理解耦为离线表示构建与在线查询执行两个阶段：离线利用 LLM 为文档生成语义表示，在线针对具体查询训练轻量级代理模型，对大部分文档进行快速筛选，仅将少量语义模糊样本交由高性能 LLM 判断。为提升筛选效率，论文提出基于对比学习的查询感知代理模型训练方法，以及满足目标精度约束的自适应级联决策机制。实验表明，ScaleDoc 在保持准确率的同时，可实现超过 2 倍端到端性能提升，并减少最高 85% 的 LLM 调用次数，大幅降低大模型语义分析成本，为大规模文档智能检索与分析提供了高效解决方案。

该研究成果论文：Hengrui Zhang, Yulong Hui, Yihao Liu, Huanchen Zhang, “ScaleDoc: Scaling LLM-based Predicates over Large Document Collections”, Sigmod 2026.

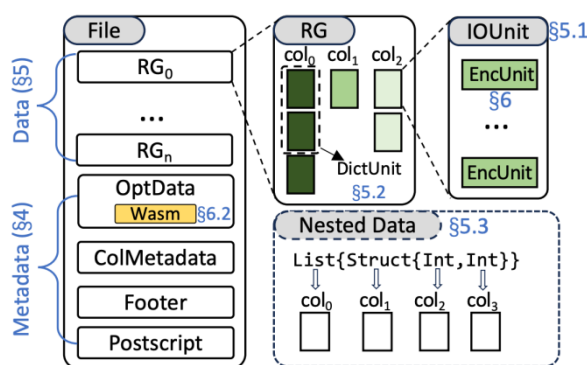


图 . ScaleDoc 总体工作流程

Group Cache: 面向 LSM-tree 范围扫描的分组缓存加速技术

团队该文针对大语言模型（LLM）在海量非结构化文档语义分析中的高推理成本问题，提出高效语义谓词执行系统 ScaleDoc。该系统创新性地将语义查询处理理解耦为离线表示构建与在线查询执行两个阶段：离线利用 LLM 为文档生成语义表示，在线针对具体查询训练轻量级代理模型，对大部分文档进行快速筛选，仅将少量语义模糊样本交由高性能 LLM 判断。为提升筛选效率，论文提出基于对比学习的查询感知代理模型训练方法，以及满足目标精度约束的自适应级联决策机制。实验表明，ScaleDoc 在保持准确率的同时，可实现超过 2 倍端到端性能提升，并减少最高 85% 的 LLM 调用次数，大幅降低大模型语义分析成本，为大规模文档智能检索与分析提供了高效解决方案。

该研究成果论文：Hengrui Zhang, Yulong Hui, Yihao Liu, Huanchen Zhang, “ScaleDoc: Scaling LLM-based Predicates over Large Document Collections”, Sigmod 2026.

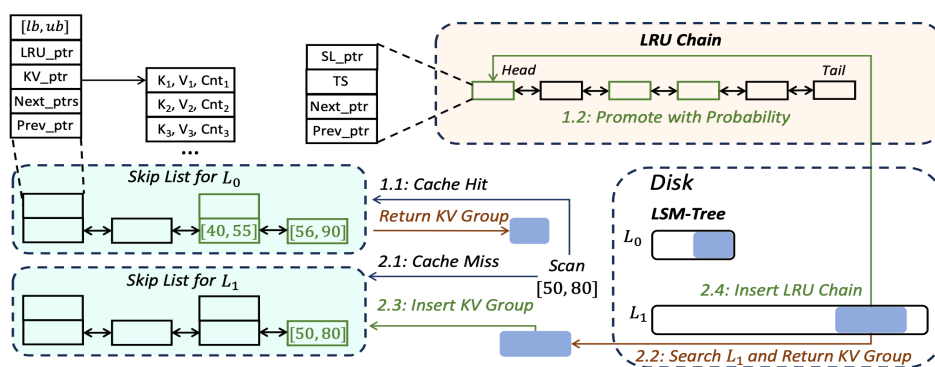


图 . Group Cache 总体工作流程

三、计算经济学

主要完成人：房智轩研究组

考虑中途状态的合约设计

合约设计是算法博弈论和多智能体系统研究中的重要问题，关注在信息不对称条件下，委托人如何通过合理的激励机制引导代理人采取期望行动。经典委托—代理模型通常假设委托人只能观察最终结果，而无法获知代理人在任务执行过程中的中间信息。然而在许多现实场景中，例如项目外包、供应链管理、众包任务和自适应任务分配，委托人往往可以在任务完成前观察到阶段性进展。这些中途状态虽然不直接产生最终收益，却可能揭示任务质量、代理人前期努力程度以及后续结果的可能性。因此，如何系统刻画并利用中间状态信息，是合约设计理论中的一个重要问题。

房智轩团队提出了一种包含多个中途状态的两阶段委托模型，用于刻画任务执行过程中逐步显现的信息。基于该模型，团队进一步提出两类新的合约机制：一种是“中途支付合约”（pay-halfway contract），允许委托人不仅根据最终结果支付报酬，也可以根据中途状态支付报酬；另一种是“中途终止合约”（terminate-halfway contract），允许委托人在观察到不利中途状态时提前终止委托过程。团队进一步分析了中途状态所包含的两类核心信息：一类是状态本身的质量，即该状态后续产生高价值结果的可能性；另一类是代理人的初始行动选择。研究表明，当中间状态不能提供额外信息时，两类新合约与标准合约效果相同；当中间状态只反映后续结果质量时，中途终止合约可以显著优于标准合约，而中途支付合约无法获得额外优势；当中间状态能够揭示代理人的初始行动时，两类新合约都可以显著优于标准合约，并有效提取中间信息带来的福利。

该研究从理论上揭示了任务执行过程中的中间信息如何影响激励机制设计，拓展了经典合约设计模型只依赖最终结果的研究范式。

该成果研究论文：Yirui Zhang and Zhixuan Fang, “The Power of Information for Intermediate States in Contract Design”, AAMAS 2026.

四、计算机安全

主要完成人：马恺声研究组

面向基于全同态加密的 AI 推理的稀疏化方法研究

许多隐私敏感的人工智能应用希望在不暴露用户原始数据的情况下完成云端推理，如医疗影像分析、金融风控和个性化智能服务等。全同态加密为这一目标提供了重要技术路径，使服务端能够直接在密文上计算，并得到与明文计算一致的结果（图 1）。然而，RNS-CKKS 等同态加密方案虽然适合支持卷积神经网络中的数值计算，却会带来远高于明文推理的计算开销。特别是用于刷新密文层级的 Bootstrapping 操作，在 ResNet 加密推理中可占总推理时间超过 50%，成为系统性能瓶颈。

马恺声团队研究了 RNS-CKKS 加密 CNN 推理中的稀疏化加速问题，提出“更少自举稀疏性”框架 Fewer Bootstrapping Sparsity (FBS)。不同于传统明文神经网络剪枝主要减少乘加运算，FBS 将优化目标转向同态加密执行中的真实瓶颈：减少 Bootstrapping 的触发频率和处理规模。该框架设计了两类稀疏模式：EMI-Sparse 通过权重稀疏化减少同态卷积所需乘法深度，从而降低 Bootstrapping 触发频率；Channel-Sparse 通过通道稀疏和图级编译优化减少 Bootstrapping 层中的密文数量，从而降低每次 Bootstrapping 的开销。团队还提出 Performance-Aware Pruning Flow，根据同态推理延迟自动选择关键层进行剪枝，在准确率与延迟之间进行迭代权衡。

团队在多个数据集上验证了 FBS 的有效性，并在 ResNet 模型上进行评估。实验结果表明，FBS 最高可实现 1.91 倍 RNS-CKKS 加密 CNN 推理加速，同时保持很小的准确率损失。该工作说明，隐私保护 AI 的高效实现不能简单沿用明文 AI 的优化目标，而需要从密码学执行模型出发重新设计算法、编译和系统优化方法。

该成果研究论文：Guiming Shi, Yuchen Wei, Zhanhong Tan, Muyang Li, Dapeng Cao, Jingwei Cai, Shuwen Deng, Kaisheng Ma, “FBS: Accelerating CNN Inference over RNS-CKKS with Fewer Bootstrapping Sparsity,” Acceptor By ACM/IEEE Design Automation Conference (DAC), 2026.

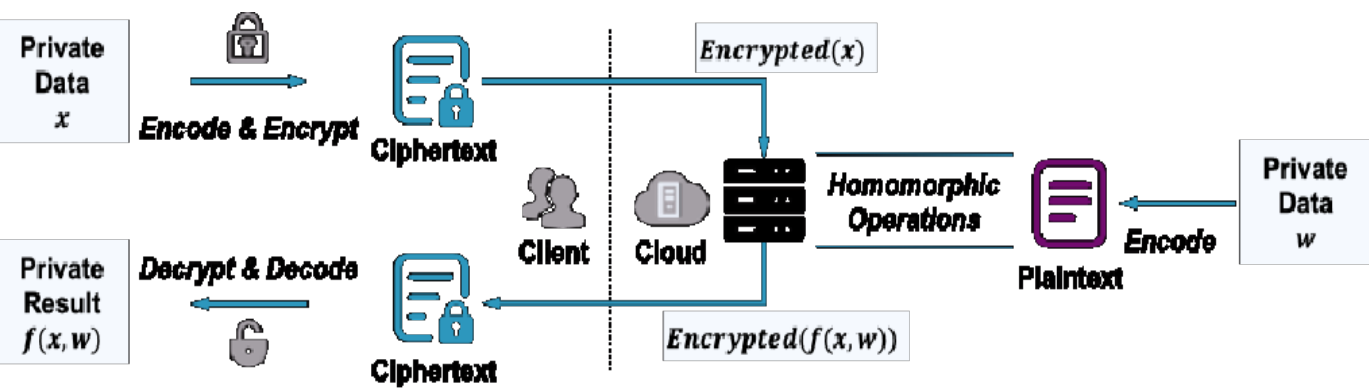


图 1 同态加密计算流程

五、密码学

主要完成人：陈一镭研究组、宋一凡研究组

交替模数学习问题、合数模数上的 Arora-Ge 算法与弱伪随机函数

在 TCC 2018 上, Boneh, Ishai, Passelègue, Sahai 和 Wu 提出了一类弱伪随机函数(weak PRFs)和强伪随机函数(strong PRFs)的候选构造, 其方法是在互素模数之间交替地求值线性函数。这类 PRF 可以由低深度电路实现, 并且对多方安全计算(MPC)十分友好。然而大家无法将这些构造归约到形式良好的数学假设, 只能直接假设这些 PRF 构造本身是安全的。

在该文中, 陈一镭团队形式化提出了一种新的假设, 称为交替模数学习问题(Learning with Alternating Moduli, LAM)。研究人员证明, 对于某些较大的模数, LAM 问题与带误差学习(LWE)问题一样困难。对于常数模数下的 LAM, 研究人员尚不知道如何将其困难性建立在 LWE 假设之上。作为替代, 研究人员给出了:

- (i) 针对素数幂模数以及某些非素数幂模数的 LAM 的多项式时间攻击;
- (ii) 通过分析经典的攻击, 为其他模数下 LAM 的亚指数级困难性提供证据。

更具体地说, 陈一镭团队提出了两种新的攻击方法。第一种攻击是一种递归算法, 用于求解具有某些模数和误差分布下的 LWE。该算法将针对素数模数 LWE 的 Arora-Ge 算法推广到了合数模数情形, 同时也能够求解某些参数下的 LAM。第二种攻击是一种多项式时间攻击, 对于任意素数 p , 它排除了在 $NC^0[p]$ 中存在弱伪随机函数的可能性。

基于上述研究, 陈一镭团队提出了若干位于 $NC^0[p_1, p_2]$ 中的弱伪随机函数候选构造, 其中 p_1, p_2 是不同的素数。这些构造的安全性基于常数模数下的 LAM 假设, 或者基于常数模数下的带舍入学习(Learning with Rounding, LWR)假设。与 Boneh 等人提出的弱伪随机函数候选方案相比, 研究人员的弱伪随机函数候选方案属于相同的复杂度类别, 同时具有一个优势, 即其安全性能够归约到形式良好的数学假设。

该成果研究论文: Yilei Chen, Liheng Ji, Wenjie Li. “Learning with Alternating Moduli, Arora-Ge over Composite Moduli, and Weak PRFs”, CRYPTO 2026.

抗 AC0 泄露的电路编译器

现代密码学系统通常被部署于复杂的现实环境中，这使得攻击者在传统的“输入 - 输出攻击”外，还可以通过对计算中间结果的侧信道攻击来窃取秘密信息。侧信道攻击对系统的安全性造成了严峻的威胁，因此研究对它们的防范成为密码学中的重要问题。弹性泄露电路（Leakage-Resilient Circuits）为侧信道攻击提供了通用的防御范式：将输入在无泄露环境下编码，在输入编码上运行原电路的无计算泄露版本获取输出编码，最后在无泄露环境下解码从而得到输出。

然而，弹性泄露电路受限于理想的无泄露编码器和解码器，在许多实际问题中发挥不尽人意。为解决这一问题，近年来研究者提出了抗泄露电路（Leakage-Tolerant Circuits）模型，该模型考虑在明文输入输出上进行直接的计算，并要求达到在此前提下最理想的安全性：任何对中间结果的侧信道攻击都可以被对输入和输出的泄露所模拟。尽管抗泄露电路从理论上给出了最强的安全性，但它的构造仍局限于一些最简单的泄露函数类，仍不能很好地刻画现实中的攻击手段。

宋一凡团队研究了对更复杂的泄露函数类构造抗泄露电路的问题。团队创造性地提出了具备更强安全性保证的强抗泄露电路模型，进而给出了首个抗 AC0 泄露的电路编译器。作为应用成果，团队还给出了首个支持高效模拟的抗奇偶性泄露的电路编译器，和首个对深度为 1 的连续 AC0 泄露具备弹性的确定性编译器。

该成果研究论文：Yaohua Ma and Yifan Song, “Leakage-Tolerant Circuits Against AC0 Leakage”, Crypto 2026.

在随机预言机模型下突破混淆电路大小的 $\Omega(|C|^\kappa)$ 屏障

混淆电路是构建安全多方计算协议的重要基础组件。长期以来，在随机预言机模型下，为了达到 $2 - \kappa$ 的统计安全误差，现有的混淆电路大小始终无法突破 $\Omega(C^\kappa)$ 比特（其中 C 为电路大小， κ 为统计安全参数）。尽管过去有大量工作致力于使用对称密码学原语来缩减混淆电路的体积，但普遍的观点认为 $\Omega(C^\kappa)$ 比特是该模型下的一个理论下界。

宋一凡研究组提出了一种全新的方案，首次在随机预言机模型下实现了 $o(C^\kappa)$ 比特的混淆电路大小。具体而言，对于大小为 C 、深度为 D 的电路，该方案在对抗者具有 T 次随机预言机查询限制的情况下，将混淆电路的大小成功降低至 $O(C \log T + D \kappa^2 \log T)$ 比特。这一成果成功打破了混淆电路大小长期存在的 $\Omega(C^\kappa)$ 理论屏障。

该方案的核心技术启发自多方混淆协议。与传统方案直接构造具有 $2 - \kappa$ 统计安全的单一混淆电路不同，研究组让混淆者“在脑海中”模拟执行一个具有 n 个虚拟参与者的多方混淆协议，并将协议输出作为最终的混淆电路。在此机制下，对抗者必须同时攻破 t （其中 $t = \Theta(\kappa)$ ）个虚拟参与者的电路，才能破坏整体安全性。因此，每个单一电路只需具备较低的安全性，从而将导线标签的长度从传统的 $\kappa + \log T$ 大幅缩减至 $O(\log T)$ 。

此外，研究组将此混淆电路方案扩展到了恶意安全的双方计算协议中。其验证仅需要混淆者打开一部分虚拟参与者获取的信息供计算者验证。这一方案实现恶意安全性的额外开销仅为半诚实协议通信复杂度的常数倍。同时，研究组运用 Fiat-Shamir 的技术进一步降低了通讯轮数。具体来说，在假设并行不经意传输和随机预言机的情况下，该协议首次实现了 $o(C^\kappa)$ 的通信复杂度，并且仅需 1 轮 OT 和 3 轮单向通信。

该成果研究论文：Junru Li and Yifan Song, “Breaking the $\Omega(|C|^\kappa)$ Barrier on Garbled Circuit Size in the Random Oracle Model.” CRYPTO 2026.

在 Minicrypt 中实现常数轮和线性通信的保证输出交付 MPC

保证输出交付 (GOD) 是安全多方计算 (MPC) 中最高级别的安全属性, 它确保即使有恶意节点试图偏离或破坏协议, 诚实节点最终也必定能获得正确的计算输出。然而, 在标准诚实多数 ($n=2t+1$) 的设定下, 构建此类协议极具挑战性。现有的常数轮 GOD MPC 协议往往需要与参与者数量呈指数级的通信复杂度或使用公钥密码学工具。而以往少数 Minicrypt 中 (即不使用公钥密码学工具) 实现了线性通信复杂度的方案, 其轮数均非恒定, 通常需要 $OD+n^2$ 。哪怕允许 OD 的轮数, 目前已知能达到的最优通讯复杂度也需要 $O(|C|n^3)$ 。

宋一凡研究组提出了一种全新的解决方案, 首次在随机预言机模型下, 构建了兼具线性通信复杂度和常数轮数的 GOD MPC 协议。针对大小为 C 、深度为 D 的电路, 该协议在 P2P 信道上的通信复杂度成功被降低至 $OCn\kappa + Dn^3\kappa^3 + WI \cdot \text{polyn}, \kappa$ 比特, 且广播信道通信复杂度与电路大小无关 (其中 WI 为输入数量, κ 为安全参数)。

在技术实现上, 该方案的核心逻辑是将近期一个通讯为 $O(C\kappa)$ 但仅具备“可中止安全”的多方混淆常数轮协议编译升级为 GOD 协议。其利用了可验证 P2P 信道保证了在底层协议中每条消息都被确定送达并承诺, 否则一个恶意参与方将被找到并剔除。这仅在原本协议的基础上引入了额外的 $O(n)$ 倍通信并维持了与电路大小无关的广播信道通信复杂度。这使得得到的协议依然只有线性的总通信复杂度。在此基础上, 该协议利用随机委员会虚拟化参与者, 使验证阶段能够同时精准剔除大批量恶意节点。同时, 方案采用具备纠错功能的秘密共享方案, 即便少部分虚拟节点通过提供错误混淆电路试图破坏计算, 评估者也能利用纠错码的抗错特性成功恢复线标签, 从而在无需重跑协议的情况下保证输出交付, 保证了协议最多只需要重跑一次, 维持了原本协议的常数轮性质。

该成果研究论文: Junru Li and Yifan Song, “Achieving Guaranteed Output Delivery MPC with Constant Rounds and Linear Communication in Minicrypt”, CRYPTO 2026.

具有线性通信量和 $O(n^5)$ 附加开销的统计安全异步 MPC 协议

安全多方计算 (Secure Multiparty Computation, MPC) 允许 n 个互不信任的参与方 (最多 t 个恶意参与方) 不泄露各自隐私输入的前提下共同完成计算, 是隐私计算领域的重要基础技术。现有高效 MPC 协议大多依赖同步网络假设或计算安全假设。然而, 在区块链、分布式金融和大规模互联网系统等现实场景中, 网络延迟具有较强的不确定性, 难以建立严格的全局同步时钟; 同时, 基于密码学假设的方案往往伴随较高的计算开销。因此, 设计在异步网络环境下、不依赖密码学假设且具有高通信效率的 MPC 协议, 是该领域的重要研究方向。

在考虑恶意敌手 (malicious adversary) 和最优异步 MPC 容错门限的条件下 ($t < n/3$), 此前最优成果 Goyal, Liu-Zhang, and Song [Crypto' 24] 虽然实现了与计算规模线性相关的通信复杂度, 但仍需 $O(n^{14})$ 的额外通信开销; 相比之下, 同步网络下的最优结果 Goyal, Song, and Zhu [Crypto' 20] 仅需 $O(n^6)$ 的额外开销 (同步结果的最优容错门限为 $t < n/2$)。如何缩小异步与同步网络之间的效率差距, 是该领域长期存在的重要挑战。

针对这一问题, 宋一凡团队提出首个额外通信开销仅为 $O(n^5)$ 的统计安全异步 MPC 协议 (假设共识原语 Agreement on a Common Set, ACS), 在保持最优容错能力的同时, 实现了与计算电路规模线性相关、与电路深度平方相关的通信复杂度。与此前最优成果 Goyal, Liu-Zhang, and Song [Crypto' 24] 相比, 该工作将额外通信开销 $O(n^{14})$ 大幅降低至 $O(n^5)$, 显著提升了异步 MPC 的实际应用潜力。

为实现这一突破, 宋一凡团队设计了两项关键技术: (1) 提出新的异步秘密共享 (Asynchronous Complete Secret Sharing, ACSS) 协议, 该协议允许在异步网络下让一个秘密分享者将其秘密通过 Shamir Secret Sharing 的形式分享给所有参与方, 并保证最终所有参与方得到其秘密份额。该技术实现与共享数量线性相关的通信复杂度, 并将额外通信开销降低至 $O(n^5)$ (此前最优结果 Ji, Li, and Song [Crypto' 24] 所需通讯复杂度为 $\Omega(n^{12})$); (2) 提出新的 Beaver Triple 生成协议, 基于 ACSS 技术, 实现与 Triple 数量线性相关的通信复杂度, 并将额外通信开销降低至 $O(n^3)$ (此前最优结果 Goyal, Liu-Zhang, and Song [Crypto' 24] 基于提出的新的 ACSS 技术仍需要 $O(n^7)$ 复杂度)。其中第二项关键技术的实现依赖共识原语 ACS, 其通讯开销与生成 Triple 数量无关。基于上述两个核心组件, 团队构建了新的高效异步 MPC 协议, 最终实现了整体协议通信复杂度的显著优化。

此外, 在考虑原语 ACS 的具体实现后, 该工作的总体期望额外开销增加至 $O((\kappa+n)n^4\log\kappa)$, 其中 κ 为安全参数。考虑当参与方数量 $n > \kappa\log\kappa$ 时, 该成功的通讯效率在期望上仍然超过了目前同步网络下最优统计安全 MPC 的结果 Goyal, Song, and Zhu [Crypto' 20]。

该成果研究论文: Xiaoyu Ji and Yifan Song, "Statistically Secure Asynchronous MPC with Linear Communication and $O(n^5)$ Additive Overhead", CRYPTO 2026.

六、 算法理论

主要完成人：段然研究组

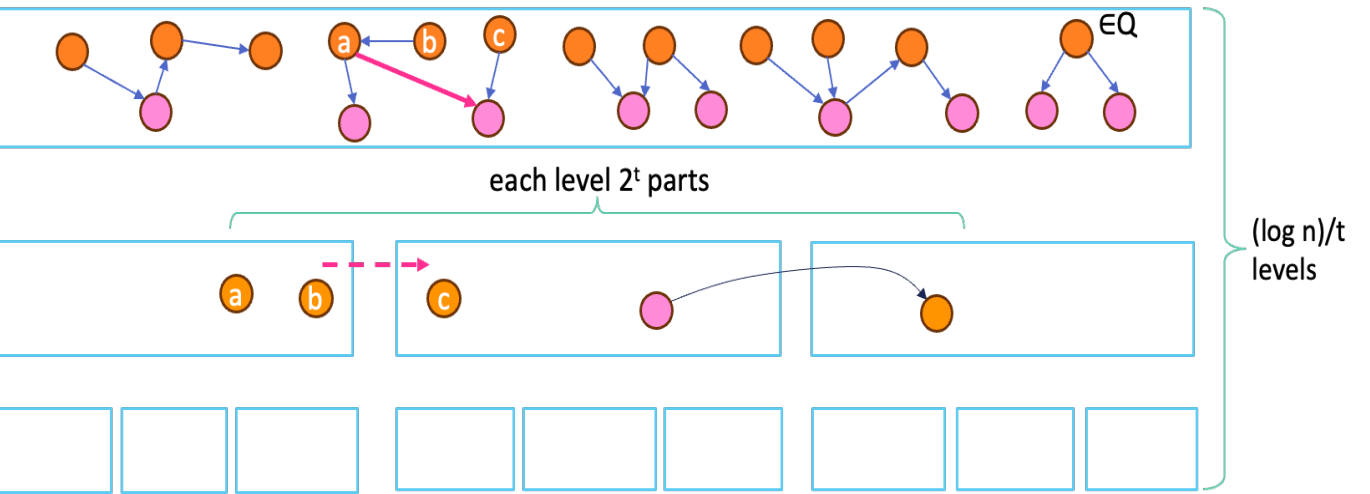
更快的有向图单源最短路算法

段然老师与研究生尹龙晖、斯坦福大学博士生毛啸、姚班校友（现德国马普所博士后）束欣凯合作的论文给出了更快的有向图单源最短路算法，时间复杂度为 $O(m\sqrt{\log n} + \sqrt{mn\log n\log\log n})$ ，在稀疏图上为 $O(m\sqrt{\log n\log\log n})$ ，改进了他们去年的复杂度为 $O(m\log^{2/3} n)$ 的算法。

最短路问题是图论领域最基础的问题之一，单源最短路问题需要找到从源点 s 到其他所有点的最短路。利用斐波那契堆的 Dijkstra 算法的时间复杂度为 $O(m+n\cdot\log n)$ ，因为 Dijkstra 算法的副产品是所有点对从 s 的距离排序，而比较模型下排序需要 $(n\cdot\log n)$ 时间，所以要改进 Dijkstra 算法就要避免整体排序。近年来关于单源最短路算法有一系列的成果，包括无向图上时间复杂度为 $O(m\sqrt{\log n\log\log n})$ 的随机算法 [Duan,Mao,Shu,Yin 2023] 和有向图上复杂度为 $O(m\log^{2/3} n)$ 的确定性算法 [Duan,Mao,Mao,Shu,Yin 2025]，以及关于 Dijkstra 算法对于距离排序问题的普遍最优性的证明 [Haeupler et al.2024]。姚班校友闫书奔去年将无向图上 $O(m\sqrt{\log n\log\log n})$ 时间的算法也做到了确定性 [Yan 2025]。

他们的新成果在有向图上将时间复杂度也改进到了确定性的 $O(m\log^{1/2+o(1)} n)$ 。新算法避免了每层单独运行 Bellman-Ford 过程，而是在每层运行局部的 Dijkstra 算法，将每层部分排序的时间均摊到边上，得到了接近单源瓶颈路和无向图单源最短路复杂度的算法。

该成果研究论文：Ran Duan, Xiao Mao, Xinkai Shu, Longhui Yin, “A Faster Directed Single-Source Shortest Path Algorithm” , To appear in the 53rd International Colloquium on Automata, Languages, and Programming (ICALP '26)





量子信息



一、量子模拟、量子传感

主要完成人：段路明研究组、邓东灵研究组、侯攀宇研究组、徐勇研究组

首次观测到多体动力学冻结现象

量子多体系统远离平衡态的演化是凝聚态物理、原子分子光学物理和量子信息科学中的核心问题之一。通常情况下，包含大量相互作用粒子的量子系统会在自身相互作用下逐渐热化，初态中的有序信息和相干特征随时间迅速消失。对于周期性驱动量子系统而言，这一问题尤为突出：外部驱动往往会持续向系统注入能量，使系统最终趋向于无特征的高温态，导致可观测的信号被抹去。这一热化过程不仅是理解非平衡量子物态的基础科学挑战，也直接限制了量子传感等应用。在许多量子精密测量平台中，自旋之间不可避免的相互作用会导致热化与退相干，进而限制有效测量时间。如何在有相互作用的多体系统中抑制热化、延长相干时间，并进一步将其转化为量子精密测量的提升，是该领域长期关注的重要问题。以往突破热化限制的方式主要依赖于引入无序产生局域化，或采用高频驱动诱导预热化。近年来，理论研究预言，在强振幅、中等频率的周期驱动下，量子多体系统可通过涌现守恒律进入“动力学冻结”状态，从而有效抑制热化。然而，在真实相互作用量子多体体系中实验观测这一现象，此前仍颇具挑战。

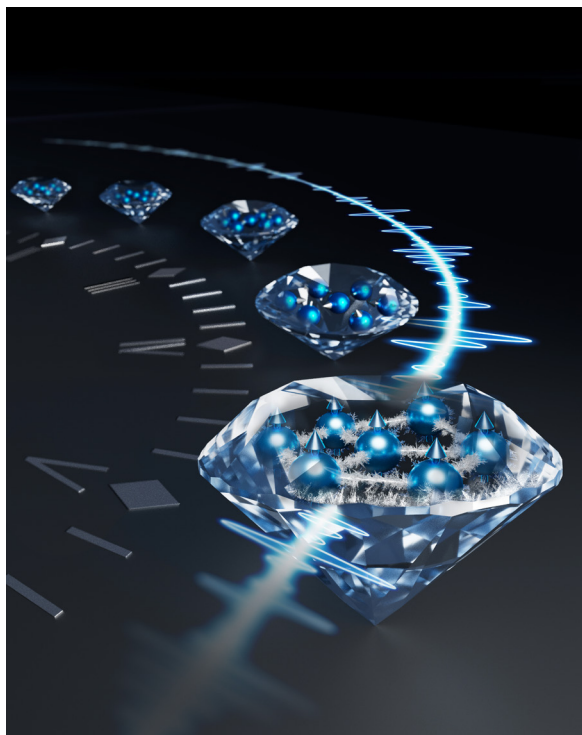


图 1：金刚石自旋系综动力学冻结示意图

针对这一挑战，研究团队利用金刚石中约一万个有相互作用的氮 - 空穴（Nitrogen-Vacancy）色心电子自旋作为实验系统，通过激光完成自旋初始化和读出，并利用全局微波场实施精确的周期驱动，使相互作用自旋系综在特定驱动参数下进入一种特殊的非平衡动力学状态。在该状态中，系统并不会快速热化，而是在长时间演化中保持集体自旋极化，表现出“被冻结”的动力学行为。实验发现，当驱动失谐与驱动频率满足特定冻结条件时，系统总自旋磁化量可在长时间内保持稳定，持续约 200 个驱动周期，超过体系相互作用限制的相干时间一个数量级以上；而当驱动参数偏离冻结条件时，系统则迅速表现出热化行为。这一结果清晰表明，观测到的长寿命磁化量保持来源于周期驱动诱导的涌现守恒律。

该涌现守恒量阻止了相互作用自旋系综的快速热化，使系统能够在远超通常相干时间的尺度上保持可观测的集体量子响应。

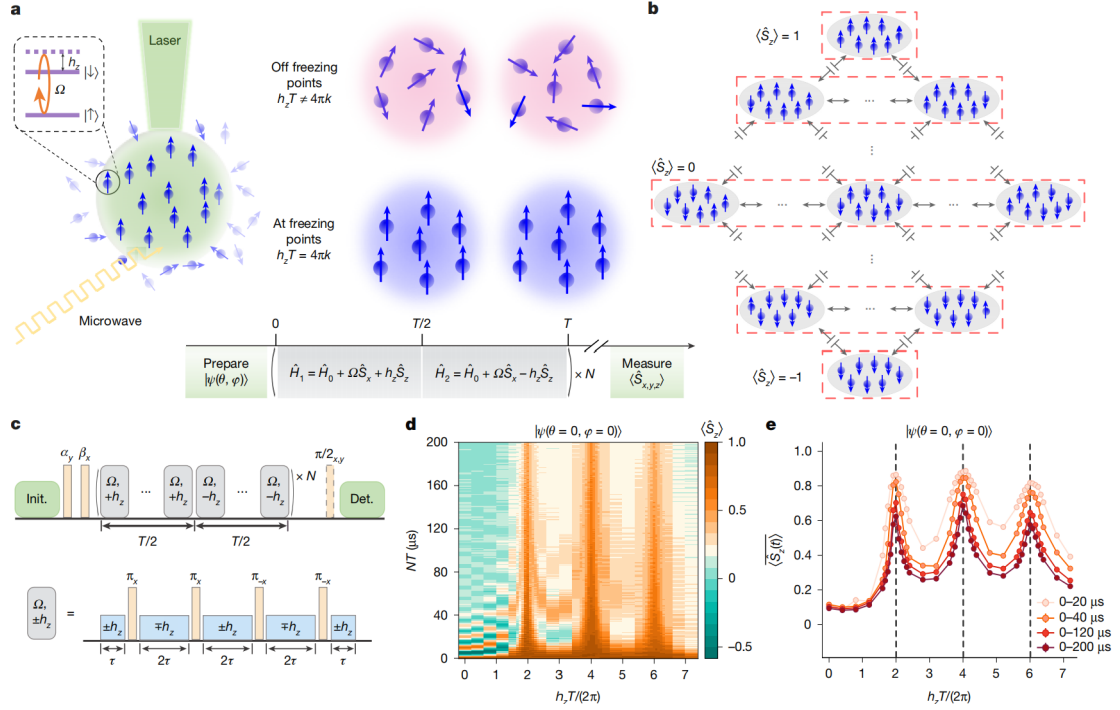


图 2：氮 - 空穴色心自旋系综中的动力学冻结现象

更进一步，研究团队将动力学冻结机制应用于交流磁场测量。传统量子传感方案通常依赖动力学解耦序列来抑制环境噪声并延长相干时间，但在相互作用较强的自旋系综中，粒子间相互作用仍会限制最佳探测时间和测量灵敏度。而该研究展示了一条不同的路径：利用多体驱动系统中的涌现守恒量来保护集体信号。实验结果表明，在相同条件下，动力学冻结传感方案相比传统周期性动力学解耦方案实现了约 2.7 倍的磁场灵敏度提升。该基于动力学冻结的量子传感方法突破了传统方案中受限于相干时间的性能瓶颈，显著增强了微弱磁信号的探测能力。

该成果不仅首次在真实固态自旋体系中观测到了多体动力学冻结的现象，揭示了一种基于涌现守恒量的新型热化抑制机制，也为发展基于多体动力学的量子传感技术开辟了新的方向，具有重要的科学意义与应用潜力。该方案仅需全局控制、易于实施，为凝聚态物理、化学及生物医学等领域中兼具高空间分辨率与高灵敏度的量子传感应用提供了切实可行的技术途径。

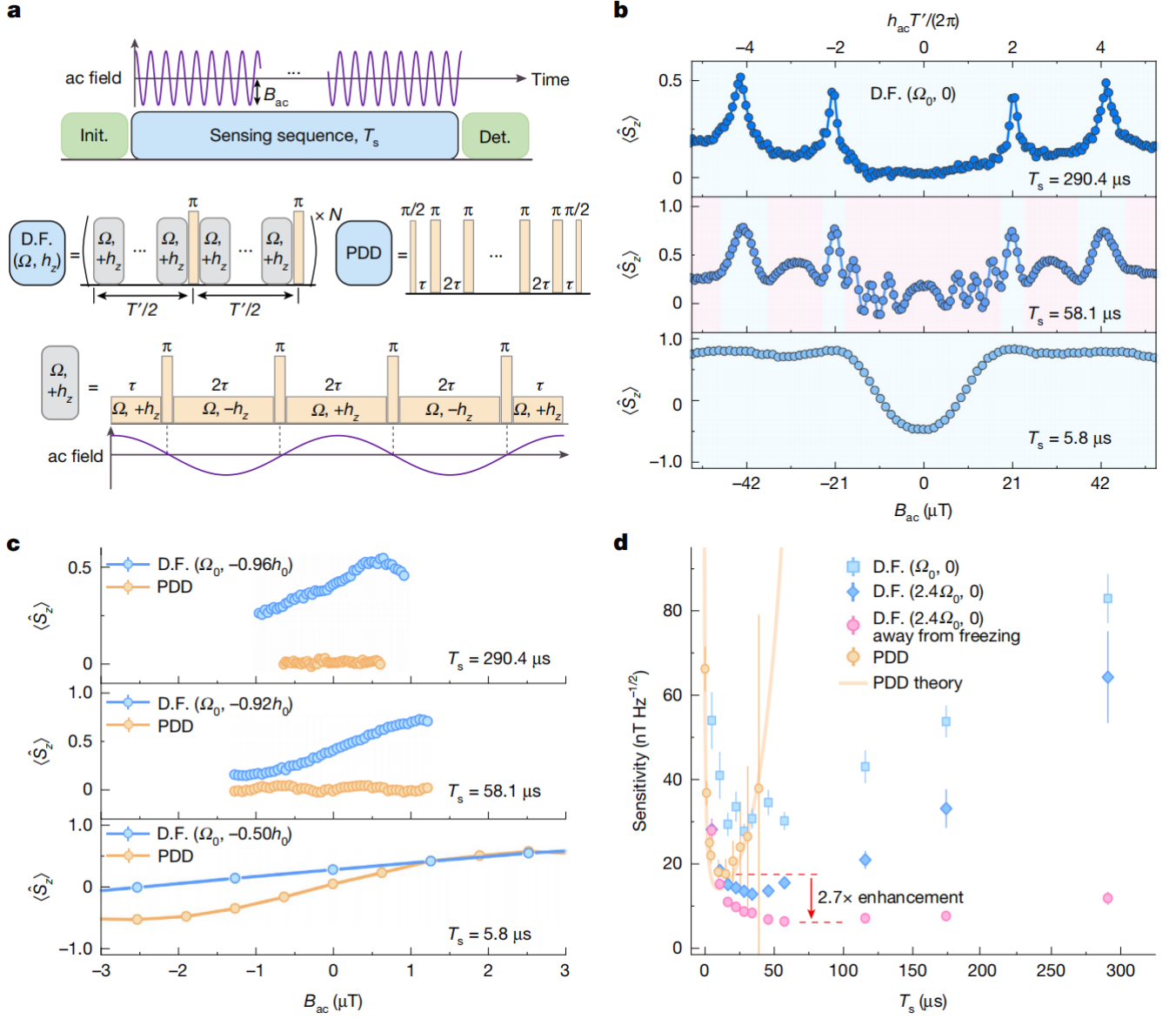


图 3：基于动力学冻结机制增强的磁场测量

该成果研究论文：Ya-Nan Lu, Dong Yuan, Yixuan Ma, Yan-Qing Liu, Si Jiang, Xiang-Qian Meng, Yi-Jie Xu, Xiu-Ying Chang, Chong Zu, Hong-Zheng Zhao, Dong-Ling Deng, Lu-Ming Duan & Pan-Yu Hou, “Dynamical freezing for magnetometry in an interacting spin ensemble”, Nature volume 653, pages1027 – 1032 (2026) .

无序里德堡原子阵列中的平均拓扑相

对称性保护拓扑相是一种新奇的物相，其稳定性依赖于特定对称性的保护。然而，真实系统不可避免地受到无序和噪声的影响，精确对称性往往被破坏，导致拓扑性质丧失。近年来的研究表明，在含噪声的系统中可以存在平均对称性，这种平均对称性能够在混合量子态中保护新的拓扑相。徐勇研究组与清华大学物理系尤力、郑盟锟团队合作，在理论上预言并首次在里德堡原子阵列中观测到了这一新奇量子态。

在里德堡阵列中，范德瓦尔斯相互作用会破坏手性对称性。尽管对于特定的某一无序构型，空间反演对称性不存在，但由给定无序强度下所有构型组成的系综仍具有平均反演对称性。这一平均反演对称性能保护拓扑相。论文提出了新的拓扑不变量来刻画该拓扑相，并发现随着无序的增强，系统经历了从拓扑平庸相到拓扑非平庸相的相变，该相变由拓扑不变量从 0 跳变为 0.5 所标志。

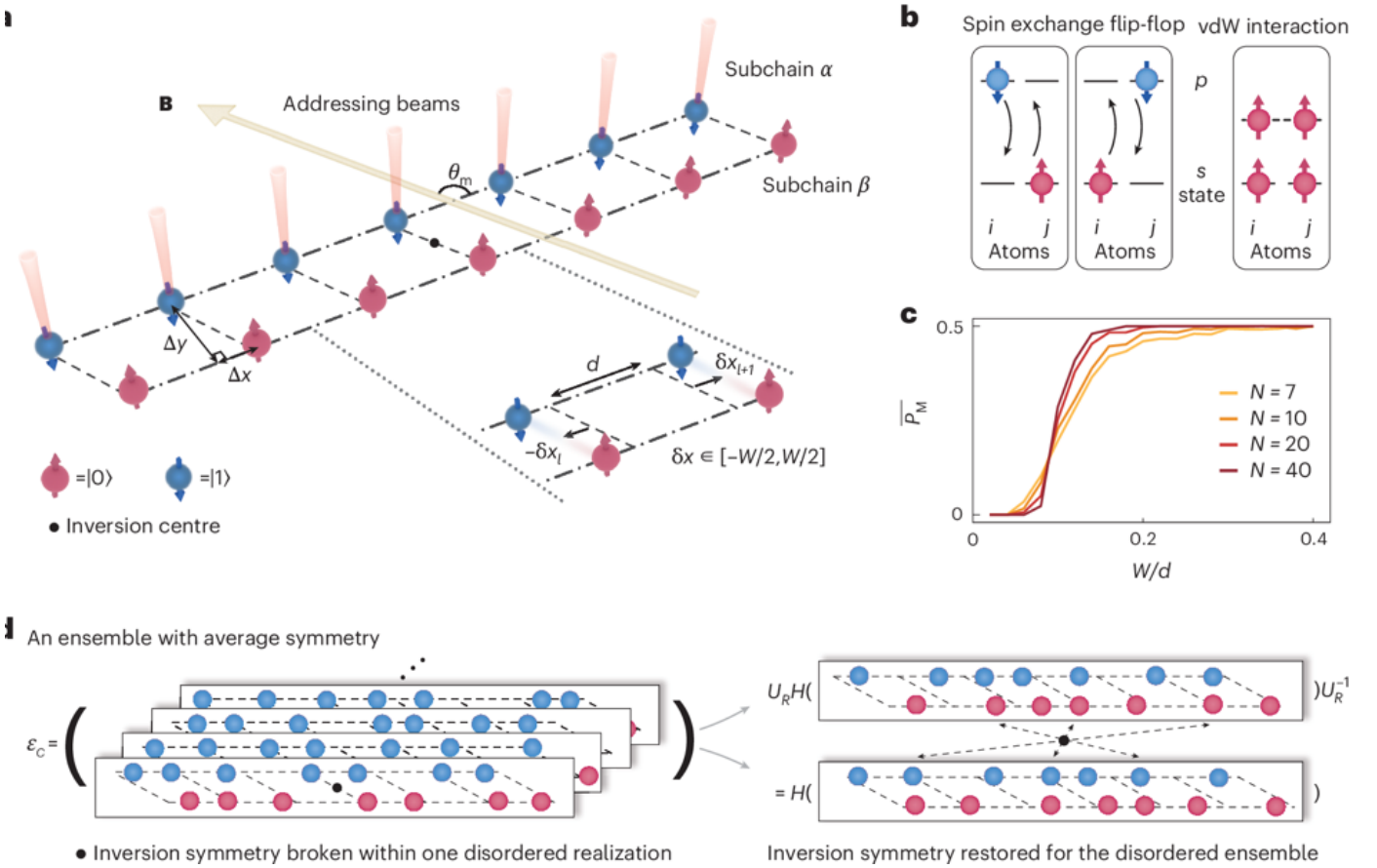


图 1: a、b, 实验中里德堡原子阵列的示意图。c, 结构无序诱导的拓扑相变。d, 平均空间反演对称性示意图。

在实验观测方面，于单粒子层面，研究团队证实了无序诱导的拓扑相变过程。当系统引入结构无序后，原本拓扑平庸的阵列涌现出近简并的边缘态，并且随着无序程度的增加，边缘格点的激发愈发显著。在多体层面，通过绝热演化制备出特定构型的基态后，无序阵列中的基态表现出显著的拓扑态特征，包括明显的边缘态以及更强的元胞间关联，与规则晶格中的平庸态形成鲜明对比。实验上还观测到了有限温度下鲁棒的边界态。

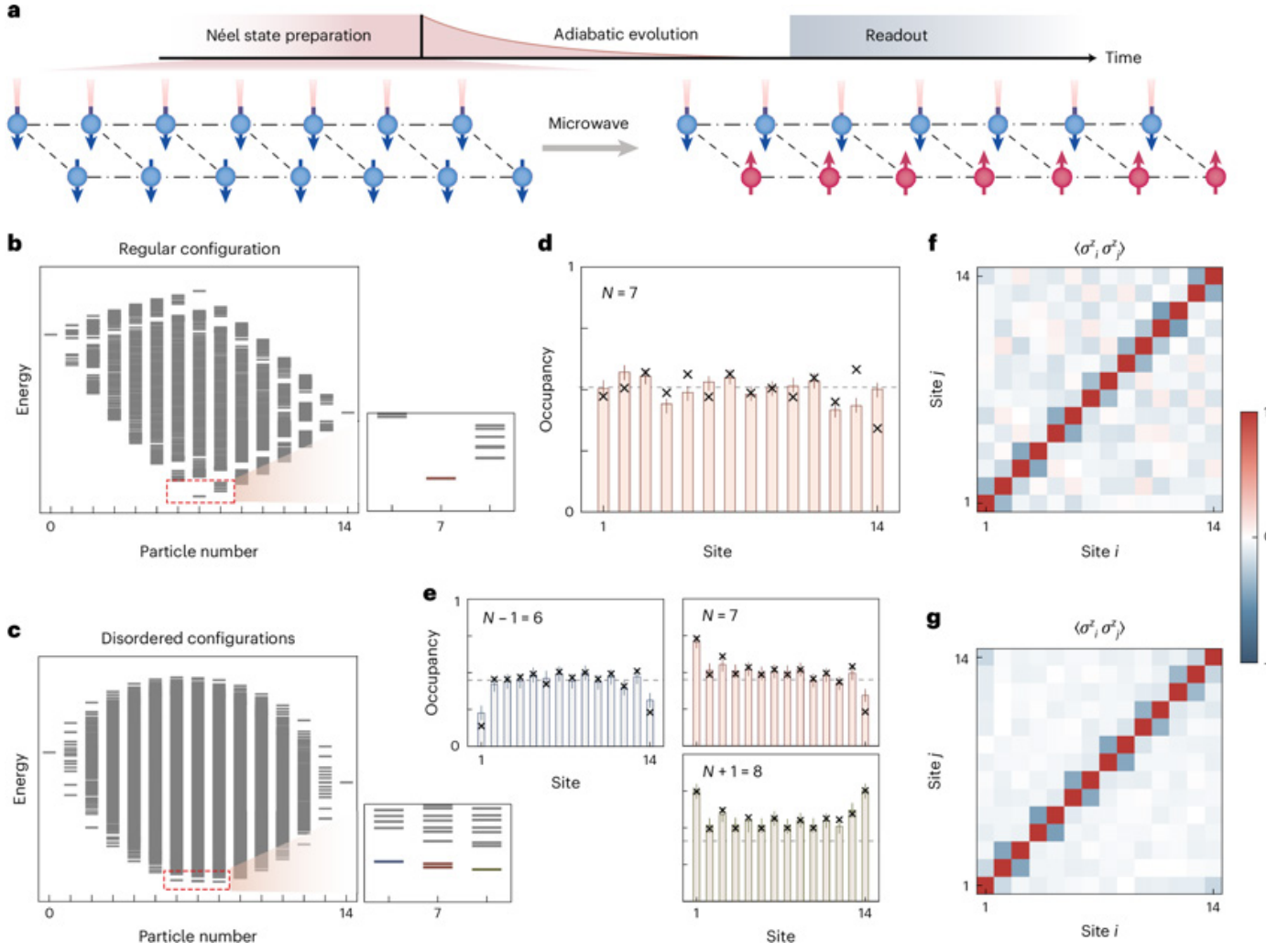


图 2: a, 实验制备多体基态示意图。规则晶格 (b,d,f) 与无序晶格 (c,e,g) 基态实验结果。

这项工作展示了里德堡原子阵列在模拟复杂多体量子问题上的巨大潜力，为在可编程量子平台上研究结构无序与拓扑的相互作用开辟了广阔的方向。

该成果研究论文：Zongpei Yue, Yu-Feng Mao, Xinhui Liang, Zhen-Xing Hua, Peiyun Ge, Yu-Xin Chao, Kai Li, Chen Jia, Meng Khoon Tey, Yong Xu & Li You, “Average topological phase in a disordered Rydberg atom array”, Nat. Phys. (2026).

二、量子信息和机器学习交叉

主要完成人：段路明研究组、马雄峰研究组、邓东灵研究组、侯攀宇研究组

机器学习识别固态自旋体系中的时间箭头

研究团队利用金刚石氮 - 空位（Nitrogen-Vacancy, NV）色心构建的十比特固态自旋量子处理器，在实验中展示了机器学习从单条量子热力学演化轨迹中识别“时间箭头”的能力。研究团队采集前向和反向量子热力学过程的测量轨迹，并利用无监督聚类、卷积神经网络和扩散生成模型，从含噪实验数据中提取由投影测量诱导的时间不可逆性。

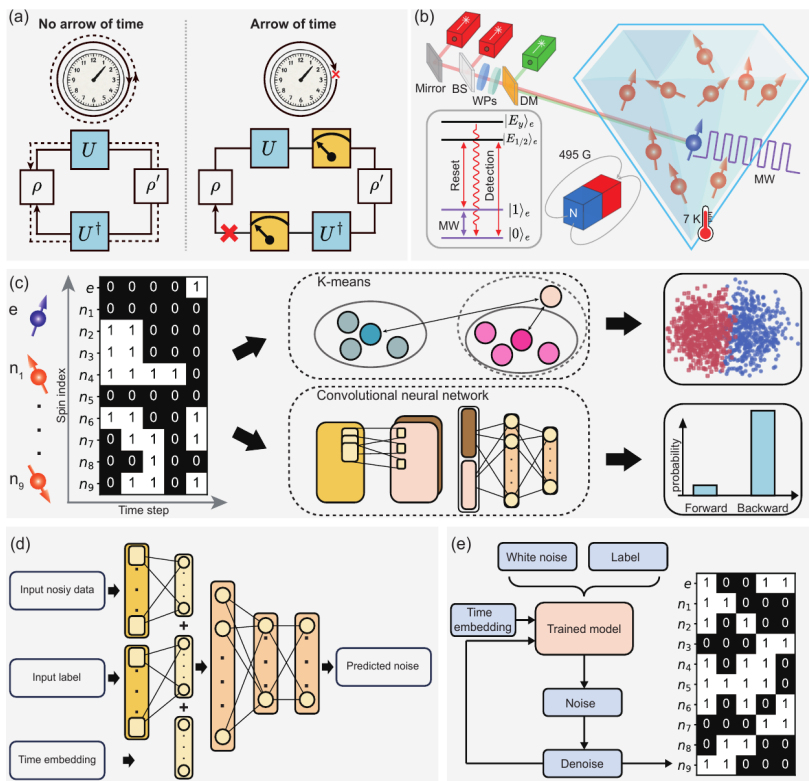


图 1：量子系统时间箭头学习框架及金刚石 NV 色心实验平台。

热力学第二定律指出，热量会自发地从高温系统流向低温系统，熵产生在平均意义下非负，这构成了宏观世界中“时间箭头”的物理基础。然而，在微观量子体系中，孤立系统的么正演化本身保持时间反演对称性，冯诺依曼熵也不随演化改变。宏观不可逆性如何从微观可逆动力学中涌现，是量子热力学长期关注的基础问题。投影测量能够通过测量反作用打破时间反演对称性并引入熵产生，但在单次微观演化轨迹中，随机涨落会掩盖这种时间方向信息，使得从真实实验数据中识别“正放”与“倒放”的演化过程极具挑战。近年来，机器学习在高维复杂数据中识别隐含模式的能力不断增强，为从量子实验轨迹中提取时间方向信息提供了新的思路。

针对这一问题，研究团队以金刚石中的单个 NV 色心为实验平台，将电子自旋作为中心量子比特，并选取其周围九个碳 -13 核自旋作为自旋环境，构成包含十个自旋量子比特的可编程量子处理器。实验在约 7 K 低温和约 495 高斯外磁场下进行，通过 532 nm 绿光制备 NV 色心负电荷态，通过 637 nm 共振激光完成电子自旋初始化和读出，并利用微波场实现电子自旋相干旋转及核自旋的通用控制。研究团队首先将中心电子自旋和核自旋环境制备在具有不同有效温

度的局域 Gibbs 热态，然后实施由交换相互作用构成的量子线路，使热量在前向过程中从较热的核自旋环境流向较冷的电子自旋。每一步么正演化之后，系统在 Pauli-Z 基下接受投影测量，测量结果共同构成一条量子热力学轨迹。相应地，团队还构造了时间反演的反向过程，并分别收集前向与反向过程各 45,000 条实验轨迹。每条轨迹可表示为一个 10 乘 5 的二值矩阵，记录十个自旋在五个时间节点上的测量结果。

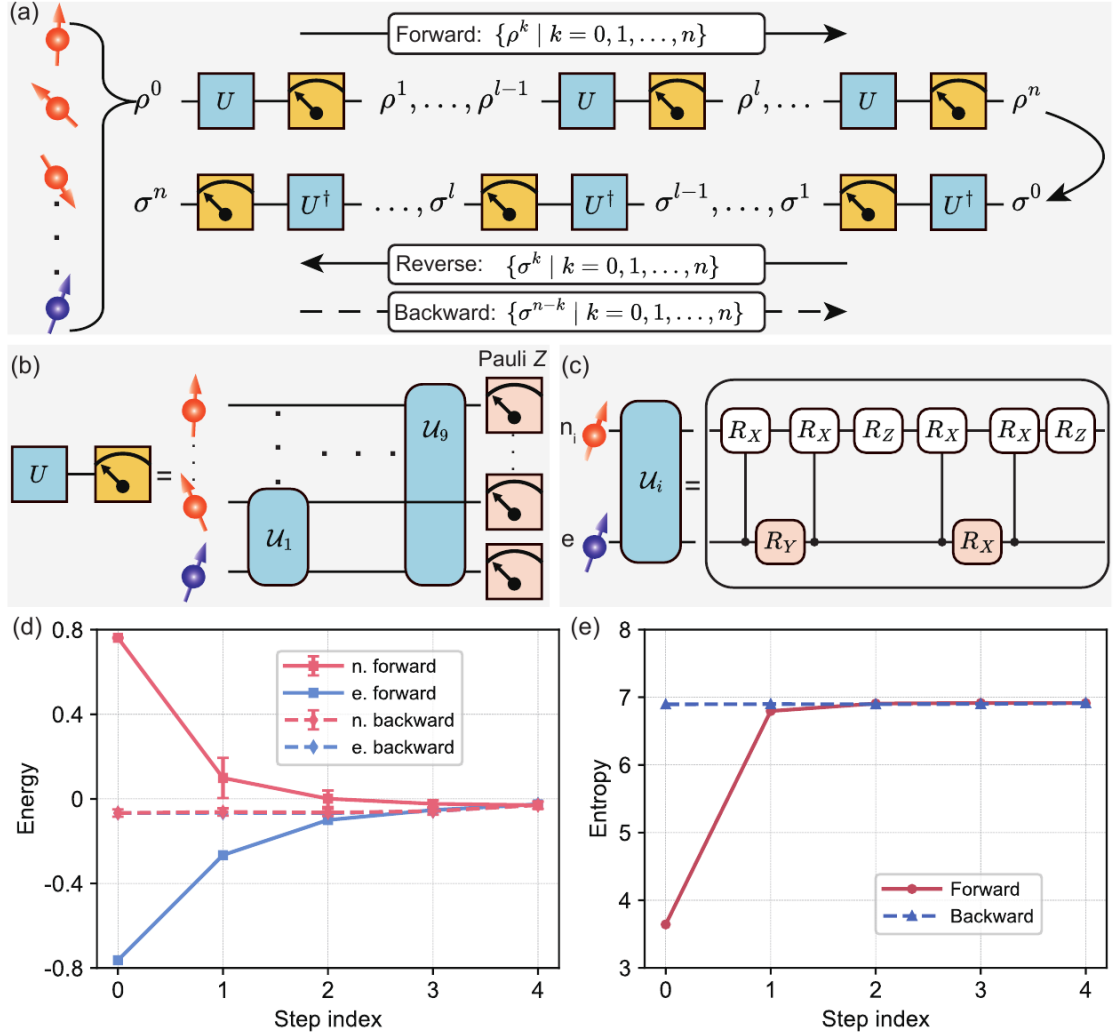


图 2: 量子热力学前向与反向过程中的能量流、熵产生及实验轨迹构造。

在数据分析中，研究团队采用两类互补的机器学习方法识别时间箭头。首先，在不向模型提供前向或反向标签的情况下，无监督 k-means 聚类算法能够将 1,000 条实验轨迹自主划分为两个簇，对应前向与反向过程，分离准确率约为 90.61%。进一步地，团队使用带标签实验数据训练卷积神经网络，其中 72,000 条轨迹用于训练、18,000 条轨迹用于测试；训练后的网络可从单条实验轨迹中判断时间方向，测试准确率达到约 92%。为了验证模型识别的并非量子线路或实验参数中的偶然特征，研究团队还构造了只包含么正演化、没有中间投影测量的数值轨迹。对于这些不存在测量坍缩的相干演化数据，同一神经网络的分类准确率退化到约 50% 的随机猜测水平，表明模型捕捉到的时间箭头来源于投影测量诱导的热力学不可逆性。

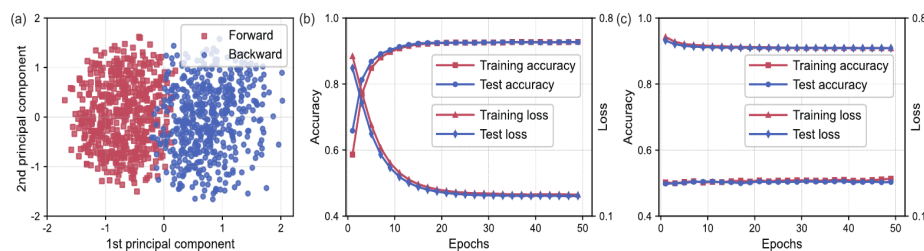


图 3：机器学习识别量子热力学时间箭头的实验结果。

更进一步，研究团队引入扩散生成模型检验机器学习是否真正学习了底层热力学过程。与判别模型只需输出类别不同，生成模型必须从噪声出发产生符合实验统计规律的轨迹。该扩散模型不预置具体物理公式，也不显式知道系统演化规律，而是通过逐步去噪生成前向或反向量子轨迹。模型训练完成后，团队生成了 90,000 条合成轨迹，并从中提取能量流动和冯诺依曼熵变化。结果显示，生成数据能够再现实验中观察到的定向热流与熵产生特征；生成态与实验投影态之间的平均保真度在训练后达到约 0.982。

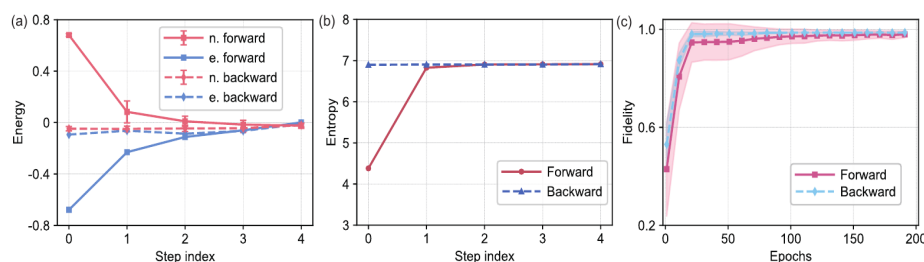


图 4：扩散生成模型重构的量子热力学轨迹。

该成果在可编程固态自旋量子处理器中展示了机器学习识别微观热力学时间箭头的实验路径，揭示了投影测量、熵产生与时间反演对称性破缺之间的联系。与传统依赖预先建立物理模型的分析方法不同，该研究表明，数据驱动模型能够在真实实验噪声和有限尺寸随机涨落中提取隐含的物理过程，并进一步生成符合热力学统计特征的新轨迹。该方法具有良好的可扩展性，有望用于表征更大规模量子系统的非平衡演化，也为研究开放量子体系中测量动力学、环境耗散与量子纠缠之间的关系提供了新工具。

该成果研究论文：Xiang-Qian Meng, Zhide Lu, Ya-Nan Lu, Xiu-Ying Chang, Yan-Qing Liu, Dong Yuan, Weikang Li, Zheng-Zhi Sun, Pei-Xin Shen, Lu-Ming Duan, Dong-Ling Deng, Pan-Yu Hou, “Machine learning the arrow of time in solid-state spins”, <https://arxiv.org/abs/2603.10344>.

量子纠错码的子空间验证

量子信息处理与量子计算在多个领域展现出潜在优势，但其实际实现受到噪声的严重制约。量子纠错是构建大规模量子计算机的关键技术，然而其往往需要引入大量物理量子比特，从而带来显著的资源开销。随着系统规模的增大，编码态保真度的验证也变得愈发困难。传统方法（如直接保真度估计和量子态验证）要么资源消耗巨大，要么仅适用于特定目标态。

为弥补这些不足，马雄峰及其博士生陈俊杰与复旦大学、芝加哥大学、香港大学合作，提出了一种通用且高效的子空间验证框架，并将其与直接保真度估计相结合，可以高效且稳健地估计制备态与其目标码子空间之间的接近程度。该框架在传统假设检验的基础上进一步发展，允许测量失败的存在，从而对测量噪声具有更强的鲁棒性。团队利用图论工具，为稳定子码和低密度奇偶校验码（LDPC 码）子空间设计了相应的验证策略。与直接保真度估计相结合后，该方法可大幅降低测量成本，并支持对一般“魔态”逻辑态的验证。

该框架和方案的提出为量子纠错码及一般魔态逻辑态提供了一种高效且可行的验证途径，有助于推动其在实际量子平台上的实现。

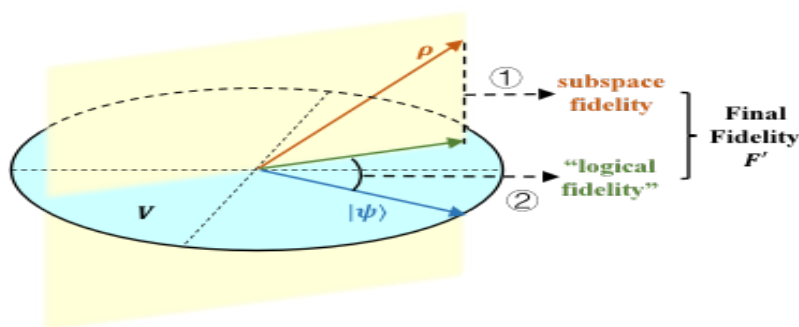


图 1 结合量子子空间验证与直接保真度估计示意图

该成果研究论文：Junjie Chen, Pei Zeng, Qi Zhao, Xiongfeng Ma and You Zhou, “Quantum Subspace Verification for Error Correction Codes”, PRX Quantum 6, 040353 – Published 3 December, 2025.

基于一维对数深度线路的近似量子纠错

高效且高性能的量子纠错是实现容错量子计算的基石。然而，构建高纠错能力的精确量子纠错码往往需要引入复杂的全局纠缠操作或高深度的量子线路，这在当前的量子硬件实现中带来了显著的资源开销和技术挑战。低深度随机线路为寻找兼顾实用性与高表现的编码策略提供了一条极具潜力的途径。然而，在受限于一维几何局部性的拓扑架构中，低深度随机线路能否以及需要达到多大的深度才能涌现出优异的纠错性能，在理论上一直缺乏严格的信息论证明与定量分析。

为解决这一问题，马雄峰及其博士生刘国定、杜振宇与合作者刘子文展开深入合作，通过严谨的信息论分析，首次证明了一维对数深度的几何局部随机 Clifford 编码线路可以展现出极其优异的近似量子纠错能力。研究表明，此类一维随机码的编码速率不仅能够达到针对泡利噪声的哈希界，还能达到针对擦除误差的信道容量。此外，研究进一步揭示了纠错后逻辑错误在超过对数深度阈值后随线路深度呈指数衰减的内在规律，从而能将译码逻辑误差降低至可忽略的水平。团队还通过证明对数深度是维持常数编码速率和高纠错性能的必要条件，确立了该类线路方案的理论最优性。为了攻克这一长久以来的证明难题，团队在技术上做出重要创新，提出并发展了一套专门适用于一维低深度线路的全新解耦定理（decoupling theorems）。

该项工作不仅在理论上完美回答了低维低深度随机线路中近似量子纠错码的涌现界限，也为未来在早期容错量子计算平台上设计低深度、高资源利用率的实用量子纠错协议与编码方案奠定了坚实的理论指导基础。

该成果研究论文：Guoding Liu, Zhenyu Du, Zi-Wen Liu, and Xiongfeng Ma, “Approximate Quantum Error Correction with 1D Log-Depth Circuits” , PRX Quantum 7, 010331 – Published 13 February, 2026.

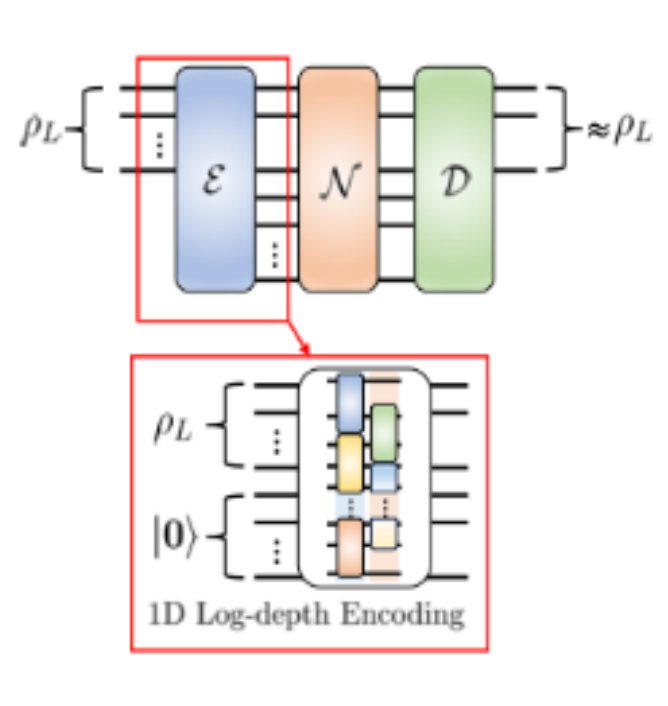


图 2 一维对数深度量子编码线路示意图

面向非线性量子性质的最优随机测量方法

高效估计量子态的非线性性质是量子学习、量子误差缓解和多体物理研究中的核心任务。然而，由于量子力学本身的线性结构，诸如态纯度、非线性期望值等二阶量子性质往往难以通过单拷贝测量高效获得。现有的经典影测量等通用方法虽然具有广泛适用性，但在估计泡利观测量、局域观测量等重要物理量的非线性期望时通常需要 $O(d)$ 个量子态样本，与已知的信息论下界之间存在显著差距。如何在不依赖量子存储和跨拷贝相干操作的条件下，构造样本复杂度最优且实验友好的随机测量协议，一直是量子学习理论中的重要开放问题。

为解决这一问题，马雄峰组博士生杜振宇、刘振寰与合作者展开深入合作，提出了一类可观测量驱动随机测量协议，用于高效估计一般形式的非线性量子性质。该协议的核心思想是将目标观测量分解为若干二值观测量，并在相应本征子空间内施加块对角随机酉演化，从而把非线性期望值的估计转化为子空间纯度之差的估计。该工作进一步面向实际量子硬件构造了高效线路实现。对于泡利观测量，团队证明 ORM 可以通过 Clifford 线路实现，同时仍保持最优的样本复杂度；对于局域哈密顿量和稳定子保真度等物理上常见的观测量，研究发展了基于泡利重要性采样的推广方案，在显著降低线路复杂度的同时保持近似最优的样本效率。此外，团队还提出了编织随机测量协议，用于同时估计多个低秩非线性观测量，在保持最优样本复杂度的同时，将所需测量基数量从经典影测量中的随系统维度增长降低到常数级，并在多观测量估计中仅引入对数级开销。

该项工作不仅从理论上回答了非线性量子性质单拷贝随机测量的最优性问题，也为量子误差缓解、虚拟冷却、混合态量子相探测和熵量估计等任务提供了更高效、更易实现的测量工具。数值模拟进一步验证了该方法在虚拟冷却任务中相较经典影测量具有显著样本优势，展示了其在近期量子计算平台和量子多体实验中的实际应用潜力。

该成果研究论文：Zhenyu Du, Yifan Tang, Andreas Elben, Ingo Roth, Jens Eisert, and Zhenhuan Liu, “Optimal Randomized Measurements for a Family of Nonlinear Quantum Properties”, PRX Quantum 7, 010360 (2026) - Published 26 March, 2026.

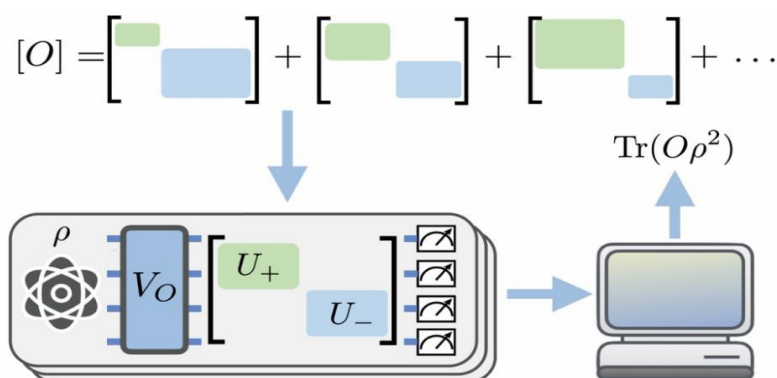


图 3 非线性量子随机测量协议示意图

仅使用局部泡利旋转实现 non-Clifford 门的基准测试

量子过程基准测试是实现高精度量子门和容错量子计算的关键前提。虽然随机化基准测试通过群旋转（Group Twirling）技术实现了对态制备与测量误差的鲁棒性，但在实际应用中，non-Clifford 门的表征一直面临巨大挑战。传统的 non-Clifford 门基准测试通常依赖复杂的多比特旋转操作，这不仅增加了理论分析和实验实现的难度，还引入了额外的旋转门噪声。

为解决这一难题，马雄峰组的叶涵、刘国定提出了一种全新的协议——泡利转移矩阵特征基准测试（PTCB）。该方案的核心特征在于，仅利用泡利操作和泡利基底下的态制备与测量，即可精确估计量子信道泡利转移矩阵（PTM）中的对称元素的乘积。研究团队通过引入虚拟的共轭 Clifford 对，在保持物理操作局域性的同时，成功访问目标信道 PTM 中的非对角元素。基于 PTCB 协议，研究团队进一步提出适用于满足 $U^2=I$ 的非 Clifford 门（如 Toffoli 门等）的保真度估计协议，并可通过特定噪声假设进一步放宽限制。以 Toffoli 门为例，数值模拟结果验证了该协议的可行性。

这项工作填补了仅使用局域操作表征通用 non-Clifford 门的理论空白，为当前量子硬件平台提供了一种高可行性的表征工具，对于大规模容错量子计算的性能评估具有重要意义。

该成果研究论文：Han Ye, Guoding Liu and Xiongfeng Ma, “Benchmarking non-Clifford gates using Pauli twirling”, Phys. Rev. A 113, 012626 – Published 30 January, 2026.

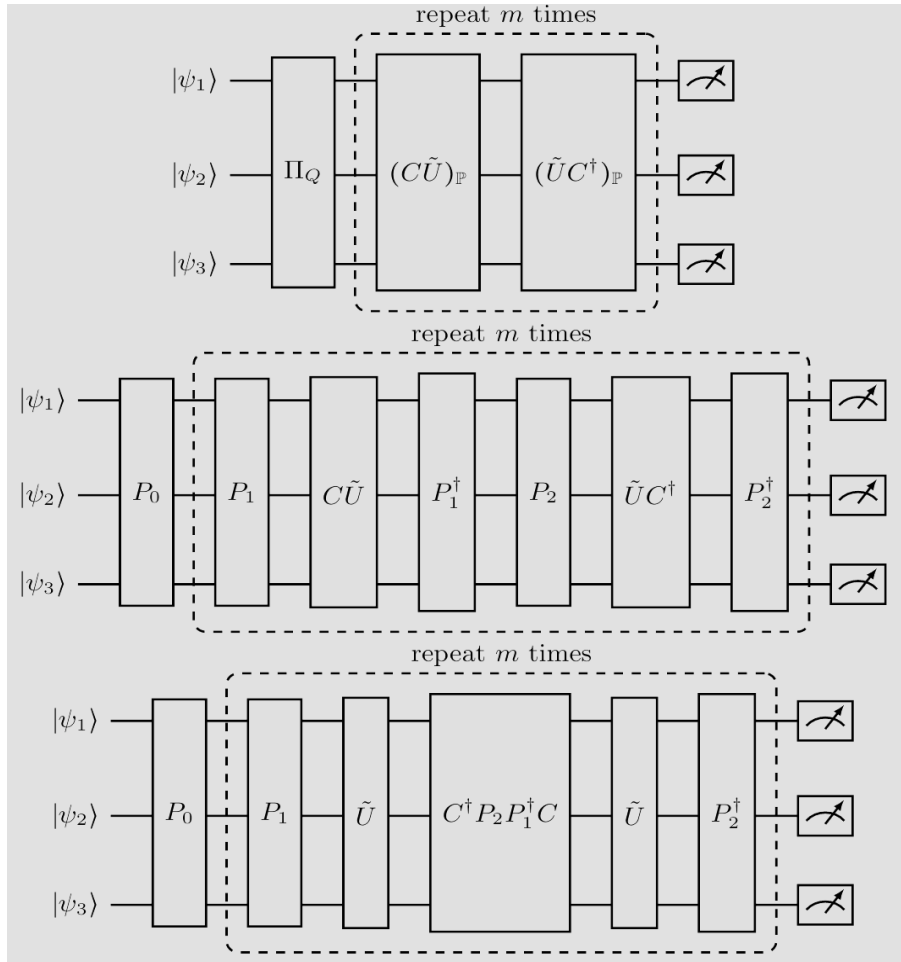


图 4 PTCB 协议的电路图及其等价形式

用于可扩展量子中继器的长寿命远程离子 - 离子纠缠

量子网络有望为安全通信、分布式量子计算和高精度量子测量提供重要支撑，而长距离纠缠分发是实现可扩展量子网络的核心前提。由于光纤传输中的指数损耗，直接分发光子纠缠难以扩展到更远距离。量子中继通过量子存储、纠缠交换和纠缠纯化等技术为突破这一限制提供了关键路径，但其实用化要求远程量子存储之间的纠缠寿命必须超过平均纠缠建立时间。此前实验受限于存储寿命、初始保真度和链路效率，远程存储纠缠往往在下一轮纠缠建立前已经显著退相干。

为解决这一问题，马雄峰与中国科学技术大学、香港大学、上海微系统所等合作者合作，基于长寿命 $^{40}\text{Ca}^{+}$ 离子量子存储、高效率低噪声电信波段频率转换接口，以及高可见度单光子纠缠协议，实现了跨长距离光纤的远程离子-离子纠缠。实验中，两个独立离子节点产生的单光子被转换至 1550 nm 电信波段，并在中间站发生干涉，由单光子探测事件预报远程离子纠缠的建立。在 10 km 光纤链路中，团队实现了 550 ± 36 ms 的存储纠缠相干时间，超过 450 ms 的平均纠缠建立时间；当前一轮 Bell 对等待后一轮纠缠建立成功时，仍可保持 0.578 ± 0.006 的期望保真度。作为应用，团队还开展了器件无关量子密钥分发实验，在 10 km 光纤链路中完成有限长安全性分析并提取秘密密钥，在 101 km 光纤链路中观测到渐近极限下的正密钥率。

该项工作突破了远程量子存储纠缠寿命短于纠缠建立时间的关键瓶颈，为多级纠缠交换、纠缠纯化和可扩展量子中继提供了重要实验模块。同时，该实验显著拓展了器件无关量子密钥分发的可实现距离，展示了长寿命远程纠缠在未来城际量子网络中的应用潜力。

该成果研究论文：Wen-Zhao Liu, Ya-Bin Zhou, Jiu-Peng Chen, Bin Wang, Ao Teng, Xiao-Wen Han, Guang-Cheng Liu, Zhi-Jiong Zhang, Yi Yang, Feng-Guang Liu, Chao-Hui Xue, Bo-Wen Yang, Jin Yang, Chao Zeng, Du-Ruo Pan, Ming-Yang Zheng, Xingjian Zhang, Shen Cao, Yi-Zheng Zhen, You Xiao, Hao Li, Lixing You, Xiongfeng Ma, Qi Zhao, Feihu Xu, Ye Wang, Yong Wan, Qiang Zhang & Jian-Wei Pan, “Long-lived remote ion – ion entanglement for scalable quantum repeaters”, Nature volume 652, pages51 – 57 (2026).

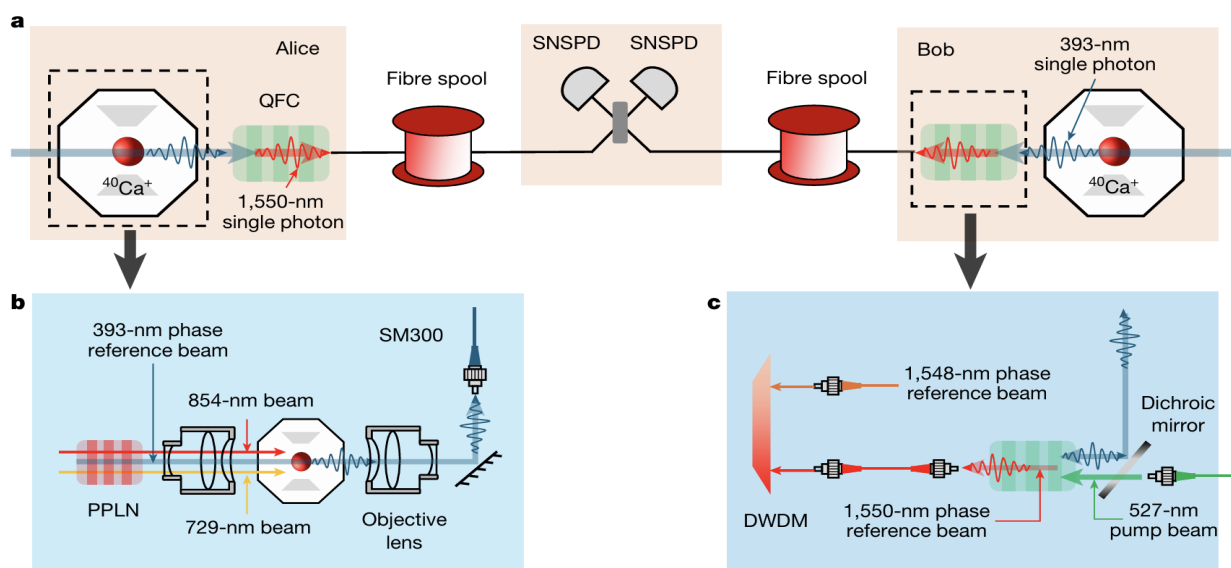


图 1 玻色编码超导量子比特研究进展示意图

三、离子阱量子网络

主要完成人：段路明研究组、濮云飞研究组

1.2 公里光纤上的多模式离子 - 离子纠缠

段路明、濮云飞研究组在离子阱量子网络方向取得重要进展。研究团队在 1.2 公里光纤链路上实现了时间多模增强的预报式离子 - 离子纠缠，并在公里级链路中保持了高保真度的远距离物质 - 物质纠缠。

量子网络的目标，是把分布在不同地点的量子节点连接成一个可以传递、存储和处理量子信息的系统。对于囚禁离子体系而言，离子量子比特具有长相干时间、高保真操控和高质量读出的优势；光子则适合在光纤中传播，承担不同节点之间的量子信息传输。因此，在远距离节点之间稳定地产生预报式纠缠，是构建离子阱量子网络的基础环节。

所谓预报式纠缠，是指两端节点先分别产生离子 - 光子纠缠，随后让两个光子在中心站发生干涉并进行 Bell 态测量。中心站的探测结果会告诉实验者哪一次尝试成功，从而把两端离子投影到纠缠态。这个方案的优点是清晰、可扩展，但在距离变长后也会遇到一个直接限制：光子飞到中心站、预报信号再返回节点，都需要时间。链路越长，系统等待的时间越多，实验重复率也越容易被通信延迟限制。

该工作关注的正是这一瓶颈。研究团队没有把远距离通信时间简单看作“空等”，而是把它重新组织成可以利用的时间窗口。在同一轮远距离通信过程中，系统连续产生多个时间模式的离子 - 光子纠缠尝试；这些光子按时间顺序进入中心站，并分别参与 Bell 态测量。只要其中某一个时间模式给出正确的预报事件，两端离子就可以被确认处于纠缠态。

实验中，研究团队搭建了两个基于钙 -40 离子的量子网络节点。两端节点发出的 866 纳米光子经光纤传输至中心测量站，并进行双光子 Bell 态测量。中心站获得的预报信号再返回节点，用来确认离子 - 离子纠缠是否建立。通过这样的结构，实验在总长 1.2 公里的光纤链路上实现了高保真度的时间多模离子 - 离子纠缠，测得纠缠保真度为 $95.9 \pm 1.5\%$ 。该保真度是目前所有物理体系中远距离（超过实验室尺度）纠缠的保真度记录。同时通过 10 个时间模式的复用获得了 4.59 倍的纠缠产生速率的增强，是基于激光冷却的物理系统中的首次实现。

该成果研究论文：Z.-B. Cui, Z.-Q. Wang, P.-Y. Liu, Y. Wang, P.-C. Lai, J.-X. Shi, Y.-D. Sun, Z.-C. Tian, H.-S. Sun, Y.-B. Liang, B.-X. Qi, Y.-Y. Huang, Z.-C. Zhou, Y.-K. Wu, Y. Xu, Y.-F. Pu and L.-M. Duan, “Temporally multimode ion-ion entanglement over 1.2km fibers,” accepted by Physical Review Letters. (2026)

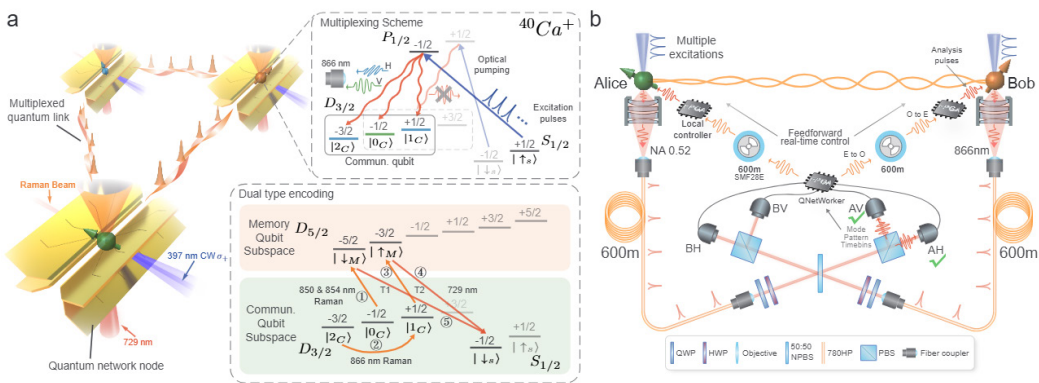


图 1：实验装置及原理示意图

四、量子计算、凝聚态物理

主要完成人：段路明研究组、吴宇恺研究组、徐勇研究组

随机横场伊辛模型中宇称限制能隙的拉伸指数标度

量子退火算法能否高效求解自旋模型的基态，其关键在于系统基态与激发态之间的能隙随体系尺寸的标度行为。对于随自旋数增加而多项式衰减的能隙，量子退火可在多项式时间内制备基态；反之，若能隙随系统规模增大呈指数衰减，则量子退火算法无法有效求解。近期有研究发现，尽管二维正方晶格上的 $\pm J$ 随机横场伊辛模型具有指数小的能隙，若将退火路径限制在宇称对称子空间内，能隙可保持多项式标度，这为量子退火解决某些自旋玻璃问题带来了希望。然而，这一结论是否适用于现实应用相关的一般随机分布尚不明确。

针对上述问题，段路明、吴宇恺研究组利用一维随机横场伊辛模型证明，只要伊辛模型的随机交换能服从独立同分布且具有非零标准差，即使在宇称对称子空间内，临界点处的能隙仍然遵循拉伸指数标度（stretched exponential scaling），并对两种典型的随机分布进行了数值验证（图 1）。研究人员进一步研究实验噪声这一常见的随机性来源，发现多项式时间量子退火算法所能容忍的实验噪声强度随系统规模增大而下降，因此大规模的量子退火应用将依赖于容错量子计算机。

该工作解析证明了宇称限制下一维随机横场伊辛模型能隙的拉伸指数标度，为理解量子退火算法的性能与适用范围提供了重要依据。

该成果研究论文：G.-X. Tang, J.-Z. Zhuang, L.-M. Duan, and Y.-K. Wu. “Stretched exponential scaling of parity-restricted energy gaps in a random transverse-field Ising model”, Phys. Rev. B 113, 144415 (2026).

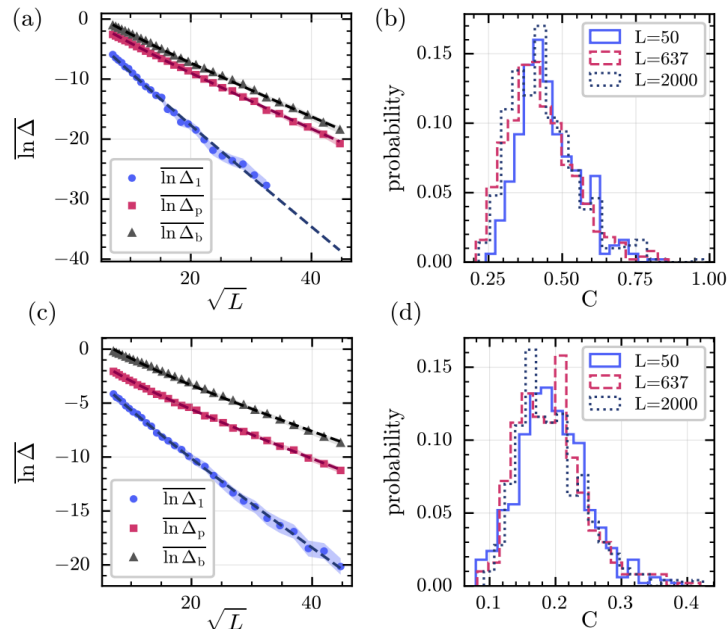


图 1：随机横场伊辛模型能隙的拉伸指数标度。

反常量子化非线性孤子泵浦

非线性效应在光学系统、玻色-爱因斯坦凝聚等物理体系中发挥着关键作用。孤子作为一种重要的非线性现象，是在传播过程中能够保持波形稳定的孤立波包，其稳定性源于非线性与色散之间的相互平衡。近期理论预测与实验观测表明，当线性哈密顿量发生周期性缓慢变化时，孤子可实现整数或分数形式的量子化泵浦。现有理论将这一孤子泵浦现象理解为孤子沿着线性哈密顿量的瞬时最大局域化单带或多带 Wannier 函数的轨迹发生移动，即孤子的量子化位移由线性布洛赫能带的拓扑陈数决定。

徐勇研究组首次在理论上发现了一种反常量子化非线性孤子泵浦，即孤子在一个周期内的位移与对应线性能带的陈数不一致。研究证实，这一反常行为源于孤子经由偶极孤子态在不同 Wannier 函数之间发生跃迁。此外，研究组还发现了非线性诱导的整数孤子泵浦——即使对应线性能带是拓扑平庸的，孤子仍能在单个周期内完成整数个原胞的泵浦。该研究成果丰富了对非线性系统中泵浦现象的认识，为非线性诱导孤子泵浦的后续相关研究开辟了全新思路。

该成果研究论文：Yu-Liang Tao, Jiong-Hao Wang, and Yong Xu, “Anomalous quantized nonlinear soliton pumping”, Nat. Commun. (2026).

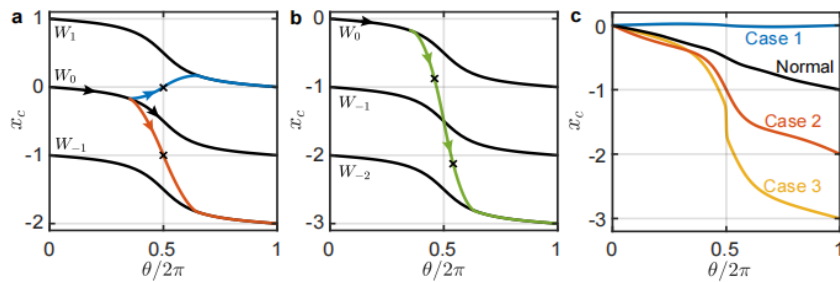


图 a-b: 反常孤子泵浦示意图; 图 c: 反常孤子泵浦质心演化图



Editor:
Kailin Li
Reviewer:
Wei Xu, Jian Li, Yipu Song