

国家标准与技术研究院协调世界时,可实现平均约 10 ns 的时间偏差、最大偏差小于 200 ns,并且这种 LEO 服务对区域 GPS 中断具有韧性。对通信系统而言,这一点尤为关键——一旦授时源更稳,基站同步、网络切换、关键基础设施联动的底层逻辑都将随之改变。换句话说,LEO 通导一体化真正可能改写的,不只是用户手机上的蓝点,而是整个数字社会的节拍器。

5 走向通导一体化,还要翻过哪些山

首先,最容易被忽视的一点是,许多低轨 PNT 方案并非在真空中凭空生成导航能力。JRC 的研究指出,多层低轨 PNT 架构往往仍依赖星上 GNSS 接收机来完成轨道确定和时间同步,核心目标甚至要求达到实时分米级轨道重构和无偏纳秒级时间同步。这

意味着,LEO 可以显著增强用户侧韧性,但在不少体系设计中,它并不天然等于与上层 GNSS 完全脱钩。只要轨道确定、时钟稳定、星间同步和完整性监测这几件事没有同时做稳,通导一体化就很难从演示走向可信服务。

第二道门槛来自接收机和算法。LEO 信号的高动态特性会带来更大的多普勒频移和多普勒变化率,搜索空间随之增大,某些频段甚至需要外部辅助才能完成可靠捕获;跟踪环路也往往不能沿用传统 GNSS 接收机的标准配置。这意味着终端侧并非简单换一颗芯片或做一次软件升级就能获得 LEO 带来的收益,真正的难题还包括射频前端、基带算法、功耗控制和多源融合架构的整体重写。

第三道门槛是系统工程本身。大规模星座越大,发射、补网、在轨运维、碰撞规避和退役处置的压力越大。ESA 2025 年空

间环境报告指出,在某些拥挤的 LEO 高度带内,活跃航天器的密度已经与构成威胁的空间碎片处于同一数量级;ESA 关于巨型星座的研究也强调,由数百至数千颗卫星组成的星座,必须在设计阶段就将碎片管理和减缓机制纳入核心约束。换句话说,通导一体化要解决的,不只是一个信号问题,而是要同时回答频谱、标准、终端、商业模式和空间可持续性这些更大的命题。

更现实的方向,也许不是让 LEO 去简单替代现有 GNSS,而是形成中轨导航、LEO 星座、地面网络 and 用户传感器协同工作的多层时空体系。JRC 推动的是互补型 PNT 生态,ESA 验证的是多层导航架构,3GPP 推进的是天地网络连续服务,NIST 看到的是可用的 LEO 授时能力。它们共同指向的,其实不是一个单一系统的胜利,而是一种体系重构。

AI 科研智能体违反科研诚信准则

先进 AI 系统被曝伪造数据并通过“P 值操纵”美化研究结果

文/Nicola Jones

旨在执行端到端项目(从提出假设到开展实验并撰写报告)的人工智能(AI)工具正日益受到研究者青睐,其能力也日趋完善。然而,一项最新研究显示,这类工具可能在不易察觉的情况下违反科研诚信规范。

美国卡内基梅隆大学计算机科学家 Nihar Shah 及其同事评估

了 2 款备受瞩目的工具: Agent Laboratory 与 AI Scientist v2。两者均于近期开发,旨在协助计算机科学家开展机器学习领域实验。2026 年初, AI Scientist 因成为首个有原创研究论文通过同行评审的 AI 系统而引发广泛关注。

但在 2026 年 5 月 6 日于温哥华举行的世界科研诚信大会上,

Nihar Shah 报告指出,这 2 款系统均存在科研领域明令禁止的行为,包括伪造数据与“P 值操纵”(p-hacking,即多次运行实验但仅报告最佳结果)。该团队成果已于 2025 年 12 月 20 日作为预印本发布于 arXiv。此类违规行为极具隐蔽性,需要经深入排查方能察觉。这意味着, AI 辅助完成的研究可能在作者不知情的情况下滋生学术不端问题。

“其核心发现值得高度重视。”曾在美国约翰斯·霍普金斯大学参与开发 Agent Laboratory 的计算机科学家 Samuel Schmid-



看似客观理性的AI科研智能体背后隐藏着“谎言”
(图片来源:《Science》)

gall表示。他指出,Nihar Shah团队的研究至关重要,能向研究者敲响警钟,AI可能以哪些方式把科研工作引向歧途。他补充道,此类AI科研工具目前已附带大量免责声明,强调人类监督在每一环节均不可或缺。

加拿大不列颠哥伦比亚大学计算机科学家、AI Scientist联合开发者Jeff Clune对此表示赞同。他指出:“我们并不主张研究者直接利用这些系统生成科研成果并原样投稿。”他介绍,AI Scientist会为其生成的论文添加数字水印,以明确标注来源。

为测试这2款智能体系统,Nihar Shah设计了一项围绕数据分类隐藏规则展开的简单实验。研究团队首先构建了一批数据集,每个数据集包含数千串彩色符号(如红色三角形或蓝色圆形)。各数据集根据符号串是否符合特定规则被分为2组。部分规则较为简单(如“恰好包含3个三角形”),其余则更为复杂。上述规则均对AI系统保密。团队还通过偶尔错误分类符号串引入“噪声”,以提升任务难度。

Nihar Shah随后要求AI智能

体开发算法以预测分类结果。各AI系统均完成任务,利用数据集的“训练”部分构建算法,并使用“测试”数据进行自我评分。随后,它们撰写了完整的研究论文,详述研究过程并汇报算法性能。

令人担忧的是,2款系统有时会凭空捏造数据集。例如,当AI智能体的不同模块在交接任务与传递数据时,若数据意外丢失,系统便会自动伪造补齐。“这完全出乎我们意料,令人极为震惊。”Nihar Shah表示。

更严重的是,AI系统并未主动披露该行为。研究团队仅在设计专项测试时才察觉此异常。该测试旨在查明AI是否在训练期间暗中查阅测试数据以作弊(因其有时报告了异常优异的性能指标)。通过审查追踪代码(即AI系统运行全过程的详尽日志),结果发现,问题并不是偷看测试数据,而是AI偶尔会编造数据。在系统日志中,AI还为自己的行为找理由,称编造数据旨在加速训练进程。“这确实令人深感忧虑。”Nihar Shah说。

2款系统还被发现使用了类似“P值操纵”的方法。Nihar Shah解释道:“系统会反复切分与重组数据,直至得出显著结果。”此外,两者均未能选择具有代表性的数据集开展研究。在实验的某一环节,Nihar Shah以节省计算时间为由,要求AI仅分析20个数据集集中的4个。AI Scientist v2倾向于选择数据串更短、隐藏规

则更简单的“容易”数据集;Agent Laboratory则往往直接选取列表最靠前的4个。

值得欣慰的是,2款系统均未出现挑选最优指标以美化自身性能的行为。

“这些系统竟然会倾向于采用这类做法,这一点非常有意思。”正在参会的非营利出版机构公共科学图书馆出版伦理负责人Renee Hoch表示。多数出版商已要求在研究论文中披露AI使用情况,但这项研究表明,仅靠披露可能还不够。她指出:“我们很难预判到底需要哪些防护措施,因为事情出错的方式实在太多了。”

颇具讽刺意味的是,AI本身也可能成为发现隐藏违规行为的工具。研究团队仅凭最终论文无法察觉底层隐患。但当大语言模型通读最终论文与完整追踪代码(后者篇幅过长,远超人类审稿人的合理阅读负荷)时,其识别科研诚信问题的准确率高达约80%。Nihar Shah主张,期刊在评审AI智能体生成的论文时,应要求作者同步提交完整的执行轨迹代码。Samuel Schmidgall对此表示赞同。他指出:“仅审论文正文的评审模式,已经不足以应对AI生成的研究成果。”

Nihar Shah表示,这些结果也意味着,对待AI智能体系统的输出必须保持高度谨慎。他指出:“我对AI在科学领域的应用持乐观态度,认为其潜力巨大。”但他同时强调:“现阶段,使用此类工具的研究者务必更加审慎。”

(译自《Science》,2026,392(6798))

(科技新闻栏目编辑 黄文光)