
高中生独立一作发论文，因参考文献造假遭撤稿

作者：writer 来源：科学网

本文原地址：<https://www.iikx.com/news/progress/41603.html>

本文仅供学习交流之用，版权归原作者所有，请勿用于商业用途！

高中生独立一作发论文，因参考文献造假遭撤稿。

编译 张晴丹

一名美国高中生Irfan Biswas去年在《医学伦理学杂志》（Journal of Medical Ethics）作为唯一作者，发表了一篇探讨人工智能在制药行业应用的论文。然而，这篇论文近日被正式撤回。调查发现，Biswas使用生成式AI“检索并解读参考文献来源”，却从未核验这些文献是否真实存在，文中引用了多篇凭空捏造的文献。

Biswas的行为并非个例。另一本伦理学期刊《学术伦理学杂志》上，一篇关于内部举报行为的论文被读者发现，29条参考文献中至少19条纯属虚构。讽刺的是，这篇论文的主题恰恰是“举报”。

据《撤稿观察》报道，近年来，虚假参考文献正在各个层面无声蔓延。而这一切指向同一个源头：以ChatGPT为代表的生成式AI。学术界正在面对的，不仅仅是几篇撤稿论文，而是一整套建立在“文献引用”之上的知识验证体系正在松动。

高中生论文撤稿，撕开AI造假的冰山一角

2025年9月，《医学伦理学杂志》刊发了一篇探讨AI在制药行业应用的论文。文章观点鲜明，认为有偏见的算法会加剧医疗体系的不公平。论文的唯一作者是美国马萨诸塞州什鲁斯伯里学校的一名高中生Biswas。

然而，这篇看似普通的论文却暗藏玄机。调查发现，Biswas使用了生成式AI来“检索并解读参考文献来源”，但投稿前并未对这些文献进行核验。结果就是，论文引用了多篇根本不存在的文献。

2026年5月28日，期刊正式发布撤稿声明。声明写道：“本刊针对该论文内容质量与参考文献准确性开展调查，其中多项疑点指向多篇参考文献为虚构。”

换句话说，论文中那些看似学术严谨的引用，完全是大语言模型凭空编造出的“学术幻觉”。更糟糕的是，调查还发现了“存在操纵同行评审的相关证据”——这意味着，这篇论文发表的背后，可能还存在更多不为人知的运作。

类似的事情，很快在另一本伦理学期刊上重演。《学术伦理学杂志》今年也做出了类似撤稿处理。一名读者发现，一篇主题为“埃塞俄比亚公立教育机构残障人士的内部举报经历”的论文中，存在大量伪造参考文献。芬兰独立研究者Erja Moore在阅读这篇论文时，本想查找几篇看起来有价值的参考文献，却意外发现“这些文献根本不存在”。

于是，Moore逐一核查了该论文全部29条参考文献，发现其中至少19条疑似人为编造。这些虚假文献存在多种情况：作者确实发表过该主题的其他论文，但此处标注的文章标题不存在；标题相近，但刊发期刊、作者完全不同；期刊卷、期、页码指向的文献与参考文献标注的内容完全不符。在所有虚构参考文献中，有11条虚构引用了施普林格·自然集团旗下的《商业伦理学杂志》。

更令人震惊的是，这种现象已经渗透进国际权威机构。2025年3月，世界银行发布了一份题为《滋养未来：依托粮食体系解决南亚肥胖问题》的报告。这份由世界银行在职工作人员参与撰写的报告，被巴基斯坦学者Muhammad Azam发现其中包含虚假参考文献。

Azam长期研究掠夺性期刊问题。他在一个运动科学研究交流群里收到群友提醒，随后仔细核查，发现多处“有问题的文献条目”。更令人担忧的是，在218份报告、政策文件、新闻类参考文献中，至少14份无法检索、无法核验真伪。

面对《撤稿观察》的问询，世界银行将这份报告从官网下架并启动核查，称“重新发布时会附上说明，清晰列明所有修改内容”。但此时，报告已被133次。

数据触目惊心：每277篇论文就有1篇伪造参考文献

这些看似孤立的案例，背后是一个系统性问题。

今年5月,《柳叶刀》刊发了一篇通讯文章,给出了令人震惊的数据:通过对PubMed收录的近250万篇论文进行核查,研究团队发现,2026年前7周刊发的论文中,约每277篇就有一篇引用了不存在的文献。生物医学文献中的伪造引用文献数量两年内增长至原来的12倍。

这一比例相较2023年的每2828篇出现一篇、2025年的每458篇出现一篇,呈现出指数级攀升态势。研究团队由美国哥伦比亚大学数据科学研究所的Maxim Topaz牵头,他们利用人工智能技术,从9710万条参考文献中识别出4406条伪造文献,分布在2810篇论文中。

更值得关注的是时间节点。Topaz团队发现,AI造假的参考文献数量在2024年年中出现激增,这与各类AI写作工具大规模普及的时间完全吻合。而根据《自然》杂志此前的报道,2025年刊发的数万篇出版物“可能包含人工智能生成的无效参考文献”。

然而,面对如此严峻的数据,学术出版界的反应却显得迟缓。核查数据显示,超过98%被检出存在虚假参考文献的论文,在调查时“未受到出版方任何处理”。

对此,不同出版机构的处理态度也存在差异。

例如,泰勒·弗朗西斯出版集团的一名发言人表示,正投入资源研发相关技术、增设专职人员并完善流程,用以筛查存在问题的引用文献。这名发言人称,若参考文献存在疑点,稿件会退回作者修改,并补充说明:“若虚假引用占比过高,或严重损害稿件整体科研可信度,稿件通常会直接拒稿。”

公共科学图书馆出版伦理负责人Renee Hoch透露,出版社正在“探索覆盖全系统的参考文献完整性筛查方案”。Hoch同时表示,公共科学图书馆不会直接将伪造参考文献定性为学术不端:“学术不端有明确界定,必须包含主观故意要素;某一问题是否构成学术不端,由作者所属机构判定,而非期刊或出版方。”

但《美国医学会杂志》前主编Howard Bauchner和《美国医学会小儿科杂志》前主编Frederick Rivara立场非常鲜明。二人主张,一旦论文出现AI生成的虚假参考文献,应当直接撤稿,因为“研究人员一旦署名成为论文作者,就要对整篇文章的全部内容承担责任”。

16岁少年的撤稿成为正面榜样

在这场由AI引发的学术诚信危机中，并非所有高中生作者的撤稿故事都令人唏嘘。2012年，一个名叫Nathan Georgette的哈佛医学院大四学生主动撤回了自己16岁时撰写的一篇论文。他发现自己第一篇论文中的一处错误假设，直接导致第二篇发表在《公共科学图书馆·综合》（PLoS ONE）的论文结论失效。

事情经过如下：2007年，年仅16岁的Georgette发表了自己第一篇学术论文，一篇免疫学相关文章刊登在《互联网流行病学杂志》。这篇论文还帮助他拿下一项知名高中奖学金。2009年，Georgette基于此前的研究成果，在PLoS ONE发表第二篇论文。

后来，Georgette进入哈佛大学读本科。2012年，他发现自己第一篇论文中存在一处错误假设。该假设对2007年那篇文章尚不构成致命问题，却直接致使PLoS ONE的这篇论文结论失效。

Georgette本可以对此秘而不宣、不了了之。但他没有回避自身失误，而是主动向PLoS ONE申请撤回该论文。当时《撤稿观察》称赞他“治学严谨、坦诚公开”，《波士顿环球报》也发表社论表示认可。

十几年后，当Georgette即将完成哈佛医学院学业、正在申请住院医师岗位时，这段经历不但没有拖累他，反而成为加分项。“当初决定撤稿时，我本以为这件事会不利于求职。”Georgette坦言，“但所有面试官都十分认可这份坦诚。”他在简历中主动标注这篇论文为“作者自主申请撤回文章”，并乐于向面试官完整讲述事情原委。这段经历让他“深刻意识到，无论临床工作还是科研工作，在所有阶段都必须开展批判性自我复盘”。

与Georgette形成鲜明对比的，是那些因AI伪造参考文献而被撤稿的作者。2025年，施普林格·自然旗下一本期刊《达鲁制药科学期刊》刊登的一篇文章，17条参考文献中有5条经核查不存在，其中一条虚构文献竟标注了《撤稿观察》联合创始人Ivan Oransky的名字——声称他撰写过一篇名为《新型监督力量正在重塑科研行业》的文章，2019年刊发于《自然》，但此文从未问世。

该期刊主编Mohammad Abdollahi，同时也是这篇论文的末位通讯作者，名下已有11篇论文遭撤稿。面对质疑，他表示，身为主编，自己“没能把控投稿审核流程”，期刊目前正在彻查事件来龙去脉。“稿件常规处理流程中出现多处疏漏，我们正在追查问题产生的过程与根源。”

Abdollahi还称，期刊配备两套独立软件，分别用于查重和检测人工智能生成内容，两套工具对这

篇稿件的检测结果均为阴性。参与稿件终稿修改的两名作者“专门用ChatGPT复核过参考文献，本想排查错误，可如今看来，这次复核毫无作用，反而催生了一系列问题”。

AI正在重塑学术研究的每一个环节，而“幻觉参考文献”只是冰山一角。当虚假信息以参考文献的形式嵌入学术数据库，当临床医生依据不存在的证据制定治疗方案，当系统综述引用的基础文献本身就是虚构——这场信任危机的影响，远比撤稿数字所显示的更为深远。

正如Topaz所说：“无论未来大语言模型能否彻底消除生成虚假内容的问题，损害已经造成。”他团队发现的4000多条伪造参考文献带来的“信息污染”，并不会随着AI技术优化而消失。对于学术界而言，真正的挑战或许不仅在于如何识别AI生成的内容，更在于如何重建读者对“参考文献”这一学术基石的基本信任。

参考链接：

<https://retractionwatch.com/2026/07/06/ethics-journal-retracts-paper-by-high-school-student-for-ai-peer-review-manipulation/#more-135257>

<https://retractionwatch.com/2025/12/02/fake-references-chatgpt-journal-academic-ethics-springer-nature-whistleblowing/>

<https://retractionwatch.com/2026/05/07/one-in-277-pubmed-indexed-papers-in-2026-shows-fabricated-references-says-analysis/>

作者：张晴丹 来源：科学网微信公众号

更多 科学进展 请访问 <https://www.iikx.com/news/progress/>

本文版权归原作者所有，请勿用于商业用途，[爱科学iikx.com](https://www.iikx.com)转发