

人工智能时代，科研诚信难题怎么解？

■本报记者 刘如楠

“科研诚信是科学事业的基石，面对人工智能带来的新形势，我们应更加坚守诚信底线，推动制度、技术与文化的协同治理，营造公开透明、诚实守信、风清气正的科研生态。”在近于北京举办的中国科学院学部科技伦理研讨会上，中国科学院学部科学道德建设委员会主任胡海岩如是说。

此次会议主题为“人工智能时代的科研诚信”，由中国科学院学部科学道德建设委员会主办。来自全国相关领域的专家学者围绕人工智能时代科研不端行为的界定和特征、审查与监管，以及人工智能在科研中的使用边界、伦理红线等内容畅所欲言，为人工智能时代推进科研诚信建设建言献策。

人工智能给科研诚信带来新挑战

自 ChatGPT 等大语言模型问世以来，人工智能已广泛应用于知识获取、数据分析、学术写作等方面，极大提升了学习和研究效率。人工智能的介入使得学术活动的主体构成发生了根本性变化。同时，“算法参与”的科研模式也导致了贡献度界定模糊、过程透明度降低、结果不可解释性增加等一系列新问题。

模型不仅“看得清”，更能“理解对” 新一代科学文献深度解析工具问世

本报讯(记者赵广立)近日,中国科学院自动化研究所(以下简称自动化所)“AI+科学”研究团队正式推出新一代科学文献解析工具——磐石·科学文献解析器,面向科研工作提供“真正‘懂科学’的智能解析引擎”。

磐石·科学文献解析器从底层算法出发,构建面向科学语义理解的多模态训练体系与强化学习机制,在公式、文本、图表等多元素协同解析上,相比传统的光学字符识别技术实现了质的飞跃。

“科学文献的识别不仅是字符的还原,更是语义结构的重建。”磐石·科学基础大模型团队成员、自动化所研究员张家俊介绍说,研究团队摒弃了仅依赖通用视觉语言大模型的思路,转而构建一套专为科学文献场景量身定制的算法训练范式。

张家俊说,该训练范式的核心在于三大技术支柱——全场景覆盖的科学数据构建、多模态监督微调策略,以及面向科学文献语义的强化学习优化机制。

在数据层面,团队采集并构建了覆盖三大典型科学书写形态的训练语料:手写体、数字排版体与纸质扫描体。手写体数据涵盖不同学者的笔迹风格、连笔习惯与轻微涂改等真实场景;数字排版体数据横跨数学、物理、天文、工程、生物、计算机等多个学科,包含大量嵌套公式、特殊符号与复杂排版;纸质扫描体数据则兼顾高清与低质量样本,模拟实际扫描或拍照中可能出现的模糊、倾斜、低分辨率等情况。所有数据均

用这种方法，二氧化碳高效变“新资源”

■本报记者 陈彬

提到二氧化碳,很多人会很自然地将其与“温室效应”“全球变暖”等联系在一起。于是,这个原本不带任何感情色彩的中性词,似乎带有了某种“贬义”。

“二氧化碳、甲醇、甲醛等物质都属于一碳化合物,这类物质的共同特点是其分子结构中均含有一个碳原子;另一个特点便是部分化合物的确可能导致温室效应。”接受《中国科学报》采访时,江南大学生物工程学院教授陈修来说。

不过,通过某些方式,这些不太“讨人喜欢”的一碳化合物,完全有可能变身造福人类的“资源”。

近日,陈修来团队成功开辟一条一碳化合物的高效转化途径,不仅能降低对化石原料的依赖,还可以将一碳化合物“变废为宝”,实现经济效益。相关成果发表于《化学工程杂志》。

费时费力的传统方式

在减少大气中一碳化合物含量、将其转化为有用资源的道路上,人们已经探索了很长时间。以二氧化碳为例,陈修来介绍,目前人们对其的资源化利用主要有两种方式。

第一种方式是通过自然界的自养微生物和植物的方式固碳。“所谓自养微生物,是指以二氧化碳为主要或唯一碳源,以无机物为氮源,通过光合作用或化能合成作用获取能量的微生物类群。”陈修来解释说,在自然界,这类微生物在碳循环中具有重要作用。比如,微藻每年可固定全球 40%以上的二氧化碳。此外,

“人工智能的应用,至少从两个方面加剧了传统的科研不端行为。一是生成式人工智能速度快,大幅降低了造假成本;二是目前大语言模型生成内容的结构、格式、语法,都呈现较高水准,检测、判断的难度急剧增加。”北京大学教授梅宏指出。

梅宏提到了最近发生的一起“操纵人工智能审稿”案例。国外某大学助理教授与他人合著的一篇论文的页面空白处,有用白色字体写的“仅限正面评价”的字样,肉眼无法看到,用鼠标光标选中后才能“露出真容”。目前,有不少审稿人借助人工智能进行审稿、给出审稿意见,如果这篇论文由人工智能进行审稿,人工智能识别出“仅限正面评价”的指令,便只会给出正面审稿意见,误导审稿人。

这并非个例。南京大学党委书记谭铁牛在报告中也提到多个操纵人工智能编造论文的案例,如编造英文词组、虚构参考文献、篡改数据等。

谭铁牛表示,人工智能至少给科研诚信带来 3 方面的新挑战,比如难以界定篡改、伪造行为,难以界定人工智能代写论文、代写代码等行为,同时人工智能参与同行评议也引发了一系列争议。

经过严格去噪、标准化标注与格式对齐,并通过均衡采样策略确保模型在多样场景下的泛化能力。

“这一全形态、多学科、高质量的数据基础,为模型理解科学表达的复杂性提供了坚实支撑。”张家俊说。

在模型训练阶段,团队采用两阶段优化策略。首先,通过多模态监督微调,使模型初步掌握文本、公式、表格、插图等异构元素的联合表征能力。在此基础上,再引入一种面向科学文献语义的梯度强化学习策略优化框架。

不同于传统以字符准确率为导向的训练目标,该强化学习策略优化框架专门设计了三重科学导向的奖励信号:公式语法正确性、符号完整性与结构合理性。模型不仅“看得清”,更能“理解对”,生成的公式在语义层面高度可靠,可直接用于符号计算、定理验证等高阶任务。

据介绍,研发团队在多个科学文献数据集上开展了系统评测,磐石·科学文献解析器在篇章级解析、公式专项识别等任务中均达到国际领先水平。

此外,为了更好满足科研需求,磐石·科学文献解析器的输出不仅包含高精度的文本与公式识别结果,还支持 JSON、Mark-down 等结构化格式输出,可无缝对接知识抽取、文献重排版、智能问答等下游应用。目前,磐石·科学文献解析器已正式开源,并作为核心组件集成于“磐石·科学基础大模型”,服务全球科研社区。

保障科研诚信关键在“人”

国家自然科学基金委员会原副主任韩宇在报告中提出,目前已经到了“人智”与“机智”共存的时代,人们需要思考,未来该以人为本还是以人机为本。

与会专家的基本共识是,在科研中,人工智能不应替代科学家进行创造性劳动和学术判断,而应作为辅助性工具。

“科研诚信问题不一定与人工智能有明确关系,在人工智能赋能科学之前,就存在各种各样的造假和不严谨问题,人工智能只是将其加速了,或者使这些问题更加隐蔽。”中国科学院自动化研究所研究员张兆翔说,“科研诚信问题与人有关,所以应该做好人的事。”

清华大学社会科学学院教授李正风表示,人工智能时代放大了传统科研诚信建设中的薄弱之处,而传统科研诚信建设的一大问题就是违规成本太低。

还有不少专家提出,要明确人工智能的使用边界和伦理红线,比如科研数据的生成、科研成果的属性和引用,必须由人来负责,科研中的价值判断、实验设计和结果解读不能完全由人工智能替代完成,科研共同体急需建立统一的人工智能使用规范和标注机制。



机器人“交警”上岗辅助指挥交通

11月3日,红灯下,机器人“交警”做出停止手势。

当日,机器人“交警”在浙江桐乡乌镇上岗,该交通指挥机器人高 1.8 米,不仅可复刻各类指挥动作,更能与信号灯系统联动,实现动态分流与智能劝导。据悉,乌镇作为世界互联网大会的永久举办地,近年来有越来越多的人工智能应用落地。

图片来源:钱晨菲 / 中新社 / 视觉中国

呼吁出台应对 AIGC 的国家政策

为应对人工智能生成内容(AIGC)在学术领域快速普及带来的挑战,防止学术不端,加强诚信治理,2023 年 9 月,中国科学技术信息研究所联合爱思唯尔、施普林格·自然、约翰威立 3 家国际出版集团发布了《学术出版中 AIGC 使用边界指南》(以下简称指南)。2025 年的最新版指南正在筹备发布中。

与会专家提到,指南目前仅作为行业推荐性标准,而非国家强制性标准。当前 AIGC 处于迅猛发展期,若成为强制性标准,可能会给整个科研生态带来较大影响,目前的时机还不合适。

北京航空航天大学人文与社会科学高等研究院副教授和鸿鹏在报告中指出,目前国内没有针对 AIGC 的整体顶层设计,高校和科研机构中也缺乏相关规范。同时,各方对于 AIGC 的使用边界还存在争议。

“我们认为,出台相关的国家政策非常必要,应该按照从无到有、先有后优的思路,明确在科研中使用 AIGC 的基本态度、指导原则、责任主体等,同时对已有政策作出调整。这样在研究机构层面,就可以依据上位政策出台具体的规范。”和鸿鹏说。

科普话强国

编者按

党的二十届四中全会提出,加快高水平科技自立自强,引领发展新质生产力。中国科协科普部特邀院士专家,畅谈科技强国。本期“科普话强国”栏目,带领读者走近天云数据雷涛团队,探寻多模态大模型如何服务于城市精细化管理和高效能运营。

“砰……”繁忙的工厂车间里,一个工人突然倒地抽搐。不过一两秒钟,警报声响起,相关工作人员迅速赶赴现场。

这并非真实场景,而是一场针对工业高风险作业的安全预演。真正的“主角”——一个不起眼的摄像头,则藏在车间上方的隐秘角落。

“城市中大量摄像头拍摄的视频,长期被视作海量、无用的数据,我们利用多模态大模型将非结构化的视频流转化成可操作、可预测的结构化知识,让每个摄像头不仅‘看得见’‘看得懂’,还能预测将会发生什么。”天云数据 CEO 雷涛告诉记者,在“智慧之眼”背后,是更强大的“智慧大脑”,守护城市安全。

2010 年,雷涛加入北京市政府支持的云计算孵化平台——云基地,开启创业之旅,但他很快发现一个问题:随着数据量越来越大,传统数据库无法满足复杂场景需求。探索新的数据处理和分析手段成了当务之急。他将目光投向分布式计算和机器学习,并在中国联通等通信运营商中进行了尝试。

“随着人工智能技术的应用,数据真正‘活’了起来。”雷涛表示,而近年来多模态大模型的发展,赋予了数据跨维度的“感知—理解—创造”能力,通过视觉—语言联合训练,将摄像头捕获的视频流内容量化为高维语义向量。

“这个过程不是简单的图像识别,而是深层理解。”雷涛进一步解释称,模型将视觉场景映射为包含行为意图、情境语义的向量表征,并匹配不同的语义空间,实现从“记录画面”到“理解语义”、从“数据堆积”到“量化知识资产”的跨越。

目前,天云数据构建的全流程人工智能安全管控系统,通过集成计算机视觉、语音识别、视频技术等多模态能力,让人工智能具备“感知—思考—行动”能力,成为高效应对各种复杂场景的“安全专家”。

“这一系统已率先在工厂中实现应用,能覆盖十大类场景中的 89 种告警规则,准确定位违规行为。”雷涛告诉记者,传统工厂的安全生产手册不仅内容繁杂,而且以文字信息为主,高度依赖工人的经验判断,无法精准、实时对工厂作业情况进行监督管理。而利用大模型自主分析安全生产手册中的内容,构建生产安全知识库,再结合多模态模型提供各类安全监察场景的实时推理分析,就能针对不同作业场景,实现基于工作流的自动规则匹配。只要出现违规行为,摄像头就能自动识别、判断,并作出警告。

“除了现场实时监控,人工智能安全管控系统还具备事前预防、事后追溯的能力,通过多维度防护,实现隐患排查和离线分析的双重保障。”雷涛表示,目前该系统针对工业安全高风险作业的覆盖率已达 89.6%,为工厂安全搭建了坚实屏障。

此外,他们正积极探索这套系统在医院、学校等场景的落地。“例如,在医院中,多模态大模型能够监控、识别并分析包括手术安全核查、个人防护装备穿戴等在内的医护作业规范,及时发现并通报人为失误和感染风险,保障患者安全。”雷涛介绍。

这一“智慧之眼”的广泛应用,正在为城市安全筑牢防线。“传统的城市管理是分割的,交通、安全、环卫、商业各管一块,缺乏横向协同,根本原因是信息孤岛导致的认知孤岛。”雷涛说,而知识化的系统整合能打通各部门边界,揭示跨域因果链。

例如,地铁故障这一交通问题会带来打车需求,刺激更多商业活动,进而加剧交通拥堵现象,导致应急车辆受阻和环境污染,影响公共安全和环境质量。“这种多米诺效应在传统体系中难以识别,而利用多模态大模型实现的知识化系统整合则使其显性化。”雷涛解释称。

这种系统整合也为网络强国建设奠定坚实基础。“网络强国建设不仅要覆盖互联网基础设施,还要涵盖电力网络、水利网络,甚至城市摄像头网络等传统基础设施,应对其进行数字化升级。”雷涛认为,这不仅是某一领域的责任,还是涉及多领域、多维度的系统性工程。

“我们必须转变思维,让多模态大模型真正服务于网络强国建设。”展望未来,雷涛认为,“目前已实现知识在不同网络间的流动,朝‘智慧大脑’甚至‘超级大脑’迈出了第一步。”

多模态大模型筑牢城市安全防线

■本报记者 赵宇彤 通讯员 徐玮彤