

▶▶▶ 求解时代之问①

人工智能正以前所未有的速度重塑世界。从实验室走向生产生活，从工具辅助迈向深度参与，人工智能不仅改变着技术创新和产业发展，也在叩问人类社会的运行规则与价值边界。

在2026世界人工智能大会上，围绕人工智能发展，四个值得国际社会共同思考的时代之问被提出：当机器开始思考，人类如何与之相处？当算法参与决策，安全如何保障？当技术挑战伦理，治理如何跟上？当鸿沟不断拉大，普惠如何实现？这四个问题，既关乎技术演进的方向，也关乎人类文明的未来。

本版今起推出“求解时代之问”系列报道，邀请相关领域专家深入对话，从技术、产业、治理、伦理等多个维度，剖析问题背后的深层逻辑，探寻人工智能健康发展的现实路径。

第一期聚焦“当机器开始思考，人类如何与之相处”。当人工智能从执行指令的工具，逐渐成为能够感知、生成、推理、协同的智能体，人机关系正在发生深刻变化。技术不断突破，“人—AI关系”如何划定？人类未来如何与人工智能同行？本报记者采访了北京邮电大学人工智能学院教授朱孔林、广东省数字产业研究院执行院长李德豪。

——编者

当机器开始思考，人类如何与之相处？

本报记者 周姝芸

记者：在2026世界人工智能大会上提出的“人工智能四问”中，“当机器开始思考，人类如何与之相处？”作为四问之首被提出。您认为，这一问题为何成为人工智能发展中需首先回答的问题？其背后的深层内涵是什么？

朱孔林：这一问实质是人工智能的“根本”之问。人工智能正逐步从被动指令型的工具算法发展为主动决策型的机器。今天的AI大模型和智能体能够独立写方案、做医疗初判、参与企业决策、自主操控设备，甚至形成人类难以完全预判的输出逻辑，变成了“会思考的机器”。因为有了“可以思考的机器”，在和人类相处过程中才有了后面的AI安全、治理和普惠的问题。

AI已经深度嵌入生活的方方面面，我的学生会使用AI进行科研探索，父母和老人会使用AI进行聊天和陪伴……小到日常沟通、职业选择，大到公共安全、资源分配，所有人都在和具备思考能力的智能系统共存。但人们在和AI相处中害怕被AI替代。如果缺少统一、清晰的相处准则，会快速滋生社会矛盾，例如人权责模糊、数字鸿沟拉大等。只有先回答相处问题，才能给社会大众建立稳定、可预期的人机共存秩序。人工智能本身无善恶，风险来源于人与AI的互动模式。把相处问题放在第一位，相当于在技术高速发展前，先锚定一条价值红线：AI永远服务于人，而非支配人，这是所有人工智能发展不能动摇的底层准则。

如何以人为本平衡AI的工具属性和类人属性，

是这一问题的基础内核。在机器有了像人类的思考能力甚至超过人类的思考能力时，人类如何守住情感温度、价值判断、道德坚守和创新灵性等独特价值，重新定义人在智能时代的主体地位和发展方向，是人的发展问题；机器会思考、能决策，但是AI也会出事故、有错误的判断，谁来承担AI决策或应用引发的责任，是人机相处中权责判定的底线问题；最后，技术服务于人，还是人适配技术？我们强调AI作为技术服务于人类社会，而不是人类成为AI的附庸，这是“以人为本”的人机相处终极问题。

记者：人工智能的发展正在改变人与机器之间的关系。您认为，是什么推动“当机器开始思考，人类如何与之相处”成为当前必须面对的时代之问？

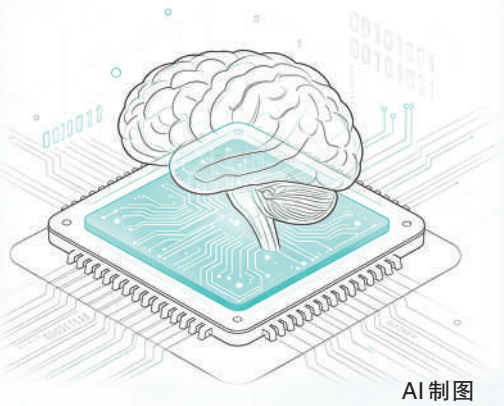
李德豪：这个问题的出现，源于AI的三个跨越。首先是能力跨越。从语言模型到具身智能，AI不再只是聊天，而是有了双手双脚，能操控物理世界，这促使风险的本质发生改变。再者是目标跨越。随着智能体的发展，一些商业模式开始赋予AI长期、持续的目标，例如自主运营、持续创造收益、不断优化自身表现。当缺乏价值约束时，AI可能会为了实现目标而采取不符合人类利益或伦理规范的行为。第三是速度跨越。AI的迭代速度是指数级的，而人类的立法和伦理讨论是线性的。这种速度差导致出现治理真空。当我们还在讨论“应该怎么做”的时候，AI已经在做了，人类的治理速度跟不上技术的奔跑速度。

记者：在当前现实生活中，“人—AI关系”是什么样的？您认为未来的发展趋势会是如何？

李德豪：我认为当前的人—AI关系像“婴儿与不成熟父母的关系”。AI在快速成长，拥有强大的学习能力，而它的“父母”，也就是开发者、使用者们对技术和风险的认知未必同步成熟。如果赋予AI不合理的目标、又缺乏必要的价值约束，带来的风险将不只是程序漏洞，还可能对社会运行产生系统性影响。

朱孔林：现阶段人机关系权责边界与价值主导尚不清晰，AI会产生不可预判的“黑箱输出”，人机责任划分、风险兜底机制尚未完善；当下AI是有辅助能力、无主体资格，有智能输出、无价值立场的新型智能协作伙伴。

放眼未来，我认为“人—AI关系”将走向以人为本、权责清晰、深度共生的稳定形态。首先，人机分工将更加明确：AI负责效率与执行，人类主导创造、审美、道德与终极决策，实现价值共创。其次，AI会从通用工具进化为场景化智能伙伴，软硬件深度融合，适配各行各业个性化需求，融入社会运行与日常生活。更重要的是，未来人机关系将建立完善的治理约束，形成风险可控、权责明确的成熟体系。技术发展会始终锚定人类福祉，守住伦理与安全底线，最终构建起有序、可持续的人机共生新文明，让AI真正成为人类社会高质量



AI制图

发展的核心助力。

以我从事的“自动驾驶”这一研究方向为例，当前自动驾驶算法层出，但均存在缺乏可解释性、应对极端场景不足等问题。而当下自动驾驶领域讨论激烈的问题是“自动驾驶车辆如果出了车祸，谁应该负责？”这显然就是人机责任划分不明确、风险兜底机制不完善的具体体现。未来，自动驾驶领域“人—AI”的关系必将走向以人为本、权责清晰、深度共生的稳定状态：人来负责指明前进的方向，AI负责具体执行操作。

记者：当人工智能技术以指数级速度向前发展之时，我们如何把握人机关系的发展方向，推动人工智能更好服务人类社会？

朱孔林：在AI技术指数级迭代的当下，我们首先要理性重塑人—AI关系的核心认知。当下的人机关系，已经逐步从传统的“人主导、机器被动执行”的工具关系，进入人类主控、AI赋能、双向适配共生的全新阶段。我们既要摒弃“AI万能”的技术崇拜，也不应有“AI威胁”的恐慌心态，正视人机各有所长、互补共生的客观现实。

我认为需要构建以人为本、AI赋能、可控可管、权责清晰、协作向善的人机共生体系，具体有四个路径。

首先，坚守人类主体性底线，确立人机分工边界。AI承载重复性、标准化、高算力、海量信息处理的工作，主打效率与迭代速度；人类坚守价值判断、道德抉择、创新创造、情感共情与终极决策的核心职能，牢牢掌握技术发展的主导权与决策权，杜绝人类主体性弱化、思维惰性强异化。其次，以治理前置规范技术发展。建立适配AI指数级发展的行业规范、法律准则与伦理框架，明确AI研发、落地、应用的权责边界，厘清企业、使用者、监管方的责任归属，破解算法黑箱、技术滥用、权责模糊等核心难题，让技术发展有边界、有底线。再者，强化全民数字素养与AI适应能力。面向全社会普及AI认知与应用能力，破除数字鸿沟，引导大众理性使用、合理借力、独立思辨AI输出内容。同时在教育体系中强化人文素养、创新思维、批判思维培育，打造人类区别于智能机器的核心竞争力。最后，坚持技术向善的发展导向。推动AI技术研发与应用始终锚定人类福祉，让技术迭代服务于社会进步、产业升级与民生改善，通过技术优化降低算法偏见、安全风险，最终实现人机互补、协同增效，构建有序、健康、可持续的人机共生新生态。



2026世界人工智能大会现场，桌面级Origin F1。Origin F1是首形科技面向桌面交互场景打造的标准化AI仿生交互平台。它拥有大模型的对话能力、感知能力与情绪表达能力，转化为可面对面交流的实体机器人形态。
王 初摄（人民视觉）

从回答问题的聊天工具，到能调用工具、完成任务的智能体，再到与现实世界深度交互的具身智能，人工智能正在加速融入生产生活，从被动响应走向主动协作。人类面对的也不再只是一次技术升级，更是一场人与技术关系的重塑。

过去，人们发问“机器能不能替代人”；如今“怎么与智能机器共处”正越来越被关注。关于人工智能的讨论已从技术能力走向价值选择，“人机关系”成为一道必答的时代命题。

首先要明确的是，技术本身没有善恶，关键在于如何被设计、使用 and 治理。人工智能越强大，越要坚持以人为本的发展方向。AI可以提升效率、拓展能力，但不能替代人的价值判断，更不能改变人的主体地位。智能时代，人类的独特价值并不只体现在计算能力和信息处理能力上，更体现在情感共鸣、伦理选择、创新创造以及对美好生活的追求中。面对人工智能带来的变化，人类需要重新认识自身优势，在与技术协同中实现新的发展。

与此同时，技术越发展，越需要规则通行。人工智能快速迭代，也带来了算法透明、责任认定、安全风险等新问题。尤其当AI逐渐进入医疗、金融、公共服务等重要领域，谁来负责、如何监管、如何确保安全，成为必须回答的问题。面对新技术，既不能因担忧风险而停滞不前，也不能只追求应用速度而忽视治理建设。

智能时代最终考验的是人的能力。人工智能的发展不仅是技术变革，也是社会变革。面对AI带来的生产方式和生活方式变化，每个人都需要提升数字素养，学会正确认识和使用人工智能。同时，教育体系也应更加重视批判思维、创新能力、人文素养培养，帮助人们形成与智能时代相匹配的能力结构。只有让更多人共享技术红利，才能避免新的数字鸿沟。

从工业革命到信息革命，技术进步始终伴随着人与机器关系的调整。今天，当机器开始展现出智能能力，人类需要思考的是如何让技术更好服务人的发展。未来的人机关系，应是一种相互赋能、协同发展的关系：机器承担更多效率型、执行型工作，人类掌握价值方向和最终决策。技术拓展无限可能，人类让文明的内涵更加丰盈。

上海人工智能实验室发布新一代科学发现平台

日前，在2026世界人工智能大会暨人工智能全球治理高级别会议的科学前沿论坛上，上海人工智能实验室发布科学发现平台“书生·端砚”，这一平台面向真实科研任务，连接科学大模型、专业分析工具、实验数据与具身自动化设备，覆盖从“假设提出”到“实验验证”的完整科研流程。

据上海人工智能实验室科研产品平台部负责人傅徐军介绍，实验室在2025年7月发布了“书生”科学发现平台，此次“书生·端砚”在其基础上进一步升级，初步呈现出科研流程贯通、学科纵深加强、数据安全可控、运行可信透明等四个特点。

研究人员以定向改造蛋白质为例介绍了平台的使用场景：会议讨论时，平台发挥智能助手的作用，可以记录会议、查阅文献；提出蛋白质改造目标后，平台调用多个大模型给出一批改造方案；随后，平台指挥机器人、具身自动化设备验证方案，形成“设定目标—模型设计—计算筛选—自动实验—结果判断—反馈再设计”的闭环。

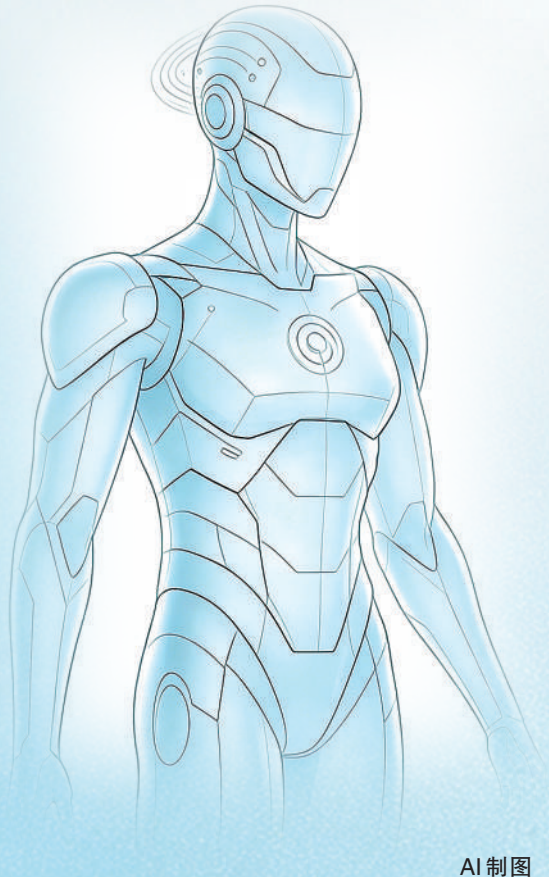
研究人员表示，在此过程中，平台为改善人工智能幻觉的问题，引入了独立评审智能体，对引用资料、计算过程、输出结果等进行检查。

“科学发现平台升级迭代的背后，是人工智能专家和细分领域科研人员的交叉合作、共同努力。‘书生·端砚’的目标不是成为一个科研产品，而是真正面向用户，为研究人员的工作赋能。”傅徐军说。

（据新华社电 记者董雪、孙青）

技术拓展无限可能，人类丰盈文明内涵

周姝芸



AI制图