

2026 世界人工智能大会
暨人工智能全球治理高级别会议
全量演讲稿汇总

华信咨询设计研究院有限公司整理

2026 年 7 月

目 录

一、主论坛（开幕式）1

（一）2026 世界人工智能大会暨人工智能全球治理高级别会议开幕式上的主旨讲话（中华人民共和国主席 习近平）1

（二）2026 世界人工智能大会暨人工智能全球治理高级别会议主席声明 7

（三）AI 发展核心拐点，从静态数据走向经验驱动（2024 图灵奖得主、“强化学习之父”、阿尔伯塔大学计算机科学教授理查德·萨顿） 12

（四）当智能体进入物理世界（阶跃星辰董事长、千里科技董事长印奇） 21

（五）要想驾驭 AI 和信息革命，战略制定要具备韧性（图灵奖获得者、美国康奈尔大学荣休教授、中国科学院外籍院士约翰·爱德华·霍普克罗夫特）28

（六）物理智能的使命是“把人还给人”“可靠运行”是行业重心（复旦大学浩清特聘教授、通用物理智能研究院院长苏昊）28

（七）科学发现的革命：让科学数据成为基础模型的原住民（中国工程院院士、之江实验室主任、阿里云创始人王坚）34

（八）无界创新与有界守护——人工智能发展的普惠与安全（商汤科技董事长兼首席执行官徐立） 42

（九）人机共生：AI 向美的实践和未来图景（欧莱雅集团全球首席人工智能官安迪雅）45

（十）AI 赋能破解生命密码——利用前沿人工智能，开发抗疾病与衰老双效疗法（英矽智能创始人兼首席执行官亚历克斯·扎沃隆

科夫)	46
二、思想者论坛	47
(一) AI 的第一性原理: 人工智能的威力及其局限 (图灵奖得主、中国科学院院士、清华大学交叉信息研究院院长及人工智能学院院长姚期智)	47
(二) 强化学习的第一性: 从经验培育超级智能 (2024 图灵奖得主、“强化学习之父”、阿尔伯塔大学计算机科学教授理查德·萨顿)	64
(三) 为机器立心: 构建 AGI 认知架构和价值体系 (北京通用人工智能研究院院长朱松纯)	78
(四) AI 未来已来: 重构商业、组织与个人 (零一万物创始人、创新工场董事长李开复)	103
(五) 探求人工智能的第一性原理: 把握硅碳时代的人机共生 (香港科技大学首席副校长, 中国工程院、英国皇家工程院、欧洲科学院院士郭毅可)	121
三、2026 人形机器人与具身智能创新发展论坛	137
(一) AI 进入物理世界: 具身智能的技术跃迁与产业拐点 (AI 战略专家、人工智能经济学奠基人、“云经济学”之父、美国未来产业研究院院长乔·韦曼 (Joe Weinman))	137
(二) 工业多臂具身智能机器人世界模型 (飒智智能董事长、CEO 张建政)	137
(三) 共生之地——上海马桥人工智能创新试验区 (上海马桥人工智能创新试验区建设发展有限公司副总经理夏吟)	138
(四) 2026 具身智能发展趋势思考 (北航机器人研究所名誉所	

长、智友·雅瑞科创平台发起人、中关村智友研究院院长王田苗)	139
(五) 具身世界模型新范式, 通往具身 ChatGPT 时刻 (银河通用机器人联合创始人、大模型负责人张直政)	139
(六) 从炫技到实干·工业具身落地探讨 (埃夫特及启智机器人董事长游玮)	140
(七) 下一脚, 踢出具身时代的安卓 (加速进化创始人、CEO 程昊)	141
四、“智序生命·AI 健未来”浦江医学人工智能论坛	142
(一) 国家卫生健康委统计信息中心党委书记、主任赵韡致辞	142
(二) 上海市卫生健康委党组书记、主任闻大翔致辞	142
五、AI 赋能民生福祉专题论坛	143
(一) 北京电控党委常委、副总经理董永生致辞	143
(二) 北电数智首席科学家窦德景	143
六、2026 中国电信人工智能生态论坛	144
(一) 星辰共生 智惠伙伴 (中国电信董事长柯瑞文)	144
(二) 专题分享 (中国电信首席技术官、首席科学家李学龙)	148
七、AI 影视专场论坛	149
(一) 通用世界模型——连接数字世界与物理世界的新基础设施 (生数科技创始人、清华大学人工智能研究院副院长朱军)	149
(二) 万兴科技创始人、董事长吴太兵	157
(三) 阿里云智能集团资深副总裁、公共云事业部总裁刘伟光	157
八、启明创投·创业与投资论坛	159
(一) 通用世界模型: 打通虚实边界, 筑基物理智能 (生数科技联合创始人、首席执行官骆怡航)	159

(二) AI 基础设施与模型能力的融合催生 (阶跃星辰联合创始人、首席技术官朱亦博)	159
九、“智见 AI 合创未来”人工智能投融资论坛	160
(一) 中国人工智能芯片需要原创技术 (国际欧亚科学院院士、清华大学教授、东方算芯董事长魏少军)	160
(二) 海通国际证券集团有限公司首席经济学家张忆东	160
(三) 国泰海通业务总监郁伟君	160
(四) 全链协同 共创 AI 黄金时代 (上海联和投资有限公司董事长、上海兆芯集成电路股份有限公司董事长叶峻)	161
十、AI 算力技术架构创新论坛	164
(一) 中国电子信息产业发展研究院党委副书记、人工智能工委会长胡国栋致辞	164
(二) 新华三集团高级副总裁、技术委员会副主席刘新民	164
十一、端侧 AI 创新和行业发展论坛	166
(一) 智能体 AI 时代：计算架构与模型效率的持续创新 (高通公司全球副总裁、中国区研发负责人、IEEE Fellow 徐皓)	166
(二) AI 手机时代来临，未来 3 至 5 年内或将迎来端侧 AGI (面壁智能 CEO 李大海)	170
十二、人工智能终端产业演进与协同发展论坛	173
(一) 中国联通副总经理郝立谦出席论坛并致辞	173
(二) 中国工程院院士、鹏城实验室主任高文，德国国家工程院院士葛兴福 (Ömer Sahin Ganiyusufoglu) 作主旨报告	173
十三、中国移动企业人工智能高质量发展论坛	174
(一) “智能驱动焕新 开放共创未来”——中国移动落实“人工智	

能+”行动创新发展实践(中国移动副总经理张冬).....	174
(二)上海东航数字科技有限公司及上海宝信软件有限公司等进行分享	176
十四、中国联通“模数同频 智造共振”AI 赋能新型工业化发展论坛	176
(一) 中国联通董事长董昕出席论坛致辞	176
(二) 元景筑原生, 万悟启工智(联通数智副总经理丁鼎).....	178
(三) AI 重构组织与运营——三一智能企业建设的新范式(三一集团 CIO 许国强)	178
十五、中国铁塔第二届“智擎未来”人工智能应用创新论坛	180
(一) 中国铁塔副总经理刘国锋出席论坛致辞	180
(二) 其他行业专家分享	180
十六、智启具身论坛 2026	182
(一)模型与数据飞轮: 冲破物理墙, 驱动物理 AI 智能涌现(觅蜂科技董事长兼 CEO 姚卯青)	182
(二) 数据驱动的具身基座模型(清华大学计算机系副研究员苏航)	183
(三) 规模化真实世界数据(Physical Intelligence 研究科学家任至意)	185
(四) 面向通用灵巧操作的人类数据与仿真扩展(Genesis AI 联合创始人王尊玄).....	185
十七、中兴通讯全栈智算生态论坛	188
(一) 中兴通讯执行副总裁谢峻石出席论坛并发言	188
(二) 中国移动研究院副院长邵春菊出席论坛并发言	189
(三) 其他嘉宾发言(清华大学教授郑纬民、中兴通讯高级副总	

裁张万春、上海仪电所属智算科技公司国产适配中心主任牛红星、阶跃星辰联合创始人、CTO 朱亦博)	189
十八、基座大模型架构创新与生态合作论坛	191
(一)什么才是真正好的人工智能产品(商汤科技董事长兼 CEO 徐立)	191
(二)“首席科学家”圆桌对话(商汤科技首席科学家林达华主持,复旦大学邱锡鹏教授,达摩院首席科学家赵德丽、阶跃星辰首席科学家张祥雨、新加坡南洋理工大学刘子纬教授)	195
十九、Enterprise AGI 价值落地与产业实证论坛	207
(一)没有企业专属数据,就没有企业专属智能(蚂蚁集团 CEO 韩歆毅)	207
(二)原国家外经贸部副部长、中国入世首席谈判代表龙永图	207
二十、共赢金砖论坛	209
(一)混合式 AI 是普惠的最优路径(联想集团副总裁陈敏仪)	209
(二)让算力像水电一样随取随用(华为数据通信产品线副总裁王辉)	209
二十一、其他论坛	211
(一)终端是 AI 普及的“最后一公里”(腾讯副总裁林松涛)	211
(二)AI 智能体要在真实场景中“干起活来”(百度创始人李彦宏)	212
(三)AI 能力正从“少数企业玩得起”走向普惠化(京东集团副总裁、京东云基础云业务负责人龚义成)	212
(四)医疗 AI 要回归智能化体系的本质(京东健康医疗 AI 负责人王国鑫)	213

（五）发展是第一位的，“不发展才是最大的不安全”（360 创始人周鸿祎）	214
（六）AI 要从“能想”到“能干”才算完成进化（网易有道高级副总裁刘韧磊）	215
（七）“技术越往前走，越要回到那个具体的个体”（智联招聘集团高级副总裁白岩）	215
（八）2026 年是“世界模型元年”（昆仑万维董事长兼 CEO 方汉）	216
（九）每个传统行业都面临底层运行逻辑重构（华为公司董事、华为云 CEO 周跃峰）	217
（十）AI 计算正从“堆卡时代”走向“优化时代”（无问芯穹联合创始人戴国浩）	218
（十一）青年一代要推动 AI 从“自动化”走向“自主化”（追觅科技创始人兼 CEO 俞浩）	218
（十二）AI 辅助诊断帮助医生在复杂罕见病例中更快锁定方向（观壹智能联合创始人、CEO 谢伟迪）	219
（十三）交易型 AI Agent 有望成为服务消费的新基础设施（滴滴集团 AI 应用负责人刘佳奇）	220
（十四）安全应从设计之初融入整车架构（曹操出行首席科学顾问、图灵奖得主约瑟夫·希发基思）	220
（十五）“技术决定能不能上路，运营决定能跑多远”（如祺出行 COO 韩锋）	221

一、主论坛（开幕式）

（一）2026 世界人工智能大会暨人工智能全球治理高级别会议开幕式上的主旨讲话（中华人民共和国主席 习近平）

演讲核心观点：

第一，坚持开放共赢，驱动创新发展。抓住难得的历史性机遇，鼓励开源开放、合作共享，全面促进人工智能科技创新、产业发展、场景应用，协同推进传统产业改造升级、新兴产业培育壮大、未来产业前瞻布局，让人工智能赋能千行百业。

第二，强化风险意识，确保安全可控。高度重视人工智能引发的各类内生和衍生风险，推动构建法律法规、技术监测、风险预警、应急响应体系，筑牢安全底线，防范滥用恶用，确保人工智能始终处于人类控制之下。共同反对在人工智能领域泛化国家安全概念、把本国安全凌驾于他国安全之上的做法。

第三，鼓励包容并蓄，促进文明互鉴。用全人类共同价值塑造人工智能的价值观，善用人工智能技术增进不同文明的理解和包容，促进不同文明交流互鉴，精心培育各美其美、美美与共的文明百花园。

第四，倡导和衷共济，完善全球治理。践行真正的多边主义，切实发挥联合国的重要作用，加强人工智能发展战略、治理规则、技术标准的对接协调，早日形成具有广泛共识的全球治理框架，让这一前沿技术更好造福人类社会。帮助全球南方

国家加强能力建设，弥合数智鸿沟，促进可持续发展，避免在人工智能领域造成新的历史不公。

来源：人民日报

<https://mp.weixin.qq.com/s/IU6cJ3THyUJyWxG1AWhJ2A>

演讲全文：

尊敬的各位同事、各位来宾，
女士们，先生们，朋友们：

70年前，一群年轻学者在美国新罕布什尔州达特茅斯会议上首次提出人工智能概念。70年间，世界各国人工智能科学家和研发者不断在未知中求索、在曲折中前行、在坚守中突破。70年后，面对新一轮蓬勃发展的的人工智能浪潮，我们相聚在浦江之畔，共同探讨促进全球人工智能朝着向上向善、造福人类的方向发展，具有重要意义。我谨代表中国政府和中国人民，对大家的到来表示热烈欢迎。

回望历史，蒸汽机的发明开启工业文明，电力的普及点亮现代社会，互联网的诞生连接整个世界。每一次科技革命都深刻重构人类生产生活方式，推动经济社会发展大幅跃升。

当前，世界百年变局加速演进，新一轮科技革命和产业变革加速突破，全球人工智能技术创新进入前所未有的活跃期，万物智联、人机协同、跨界融合、共创分享的智能技术，叠加释放巨大能量，既蕴含巨大机遇，也面临治理挑战。人类不得不直面时代之问：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？这些问题需要国际社会深入思考、共同作答。

中方认为，各国应当秉持以人为本、向上向善理念，让人工智能成为促进共同繁荣、维护共同安全的一个重要动力源，携手构建公正合理的全球人工智能治理体系。在此，我提4点意见。

第一，坚持开放共赢，驱动创新发展。人工智能是世界经济增长的新引擎和新旧动能转换的加速器，正从“数字世界”走向“物理世界”。要抓住难得的历史性机遇，鼓励开源开放、合作共享，全面促进人工智能科技创新、产业发展、场景应用，协同推进传统产业改造升级、新兴产业培育壮大、未来产业前瞻布局，让人工智能赋能千行百业。

第二，强化风险意识，确保安全可控。人工智能应当成为人类的可信工具。要高度重视人工智能引发的各类内生和衍生风险，推动构建法律法规、技术监测、风险预警、应急响应体系，筑牢安全底线，防范滥用恶用，确保人工智能始终处于人类控制之下。同时，应当共同反对在人工智能领域泛化国家安全概念、把本国安全凌驾于他国安全之上的做法。

第三，鼓励包容并蓄，促进文明互鉴。人工智能发展和应用不应该侵蚀和损害世界文明多样性和各国文化独特性。要用全人类共同价值塑造人工智能的价值观，善用人工智能技术增

进不同文明的理解和包容，促进不同文明交流互鉴，精心培育各美其美、美美与共的文明百花园。

第四，倡导和衷共济，完善全球治理。人工智能是全人类共同的智慧结晶和宝贵财富。要践行真正的多边主义，切实发挥联合国的重要作用，加强人工智能发展战略、治理规则、技术标准的对接协调，早日形成具有广泛共识的全球治理框架，让这一前沿技术更好造福人类社会。要广泛开展国际合作，帮助全球南方国家加强能力建设，弥合数智鸿沟，促进可持续发展，避免在人工智能领域造成新的历史不公。

女士们、先生们、朋友们！

今年是中国“十五五”开局之年。“十五五”规划为未来5年中国经济社会发展擘画蓝图，也为国际社会提供机遇清单。近年来，中国积极拥抱人工智能发展，坚持有效市场和有为政府相结合，加强人工智能科技创新，积极推进“人工智能+”行动，培育各类主体共生共荣的健康生态，智能经济核心产业规模已经超过万亿元人民币。各类智能终端走进千家万户，切实惠及民生，“中国智造”已成为中国式现代化的又一亮丽名片。

同时，中国始终坚持发展和安全并重，把握人工智能发展趋势和规律，不断完善相关法律法规、政策制度、应用规范、伦理准则，确保人工智能安全、可靠、可控，让人工智能这匹千里马既跑得快又跑得稳。

作为负责任大国，中国在人工智能领域始终致力于做国际公共产品的提供者。我提出《全球人工智能治理倡议》以来，中方先后推动联合国大会协商一致通过人工智能能力建设国际合作决议，发布《人工智能能力建设普惠计划》和《“人工智能+”国际合作倡议》，倡议成立世界人工智能合作组织，源源不断贡献中国方案。

中国人常讲，单弦不成音，独木不成林。人工智能发展不应该是某个国家的独奏，而应当是全球合作的交响。在各方共同努力下，世界人工智能合作组织在上海应运而生，一年前的构想已经转化为实际成果。这是中方响应全球南方呼声、团结国际社会积极推动人工智能发展和治理的重大举措，将成为人工智能发展史上的一个重要里程碑。

为进一步支持全球人工智能发展、推进全球人工智能能力建设，我宣布，未来5年，中国将面向发展中国家提供5000个人工智能专题研修培训名额；面向东盟、阿盟、非盟、拉共体、上合组织、金砖国家建设国际人工智能应用合作中心；推动气象智能预警方案“妈祖”在30个国家落地应用，守望万家灯火，护佑四海安澜。

女士们、先生们、朋友们！

“明者因时而变，知者随事而制。”人工智能技术发展越是一日千里，向上向善、造福人类的方向越要正确锚定，监管治理的尺度越要精准把握，防范失控的措施越要及时完善。要始

终用人类智慧和国际共识引领人工智能发展，使其真正成为增进全人类福祉、推动人类文明进步的强大正能量。

中国愿以更加开放的姿态、更加务实的行动、更加长远的目光，同各方一道把握和应对人工智能发展的机遇和挑战，携手共创人类社会更加美好的未来！

谢谢大家！

来源： 人民政府网

[https://www.gov.cn/yaowen/liebiao/202607/content_7075874.
htm](https://www.gov.cn/yaowen/liebiao/202607/content_7075874.htm)

(二)2026 世界人工智能大会暨人工智能全球治理高级别会议主席声明

全文：

2026 年 7 月 17 日至 20 日，2026 世界人工智能大会暨人工智能全球治理高级别会议在中国上海举行。中华人民共和国主席习近平出席大会开幕式并发表主旨讲话。哈萨克斯坦总统托卡耶夫、柬埔寨首相洪玛奈、泰国总理阿努廷、联合国秘书长古特雷斯，以及来自 100 多个国家和国际组织的官方代表、产学研界嘉宾等出席大会。在新一轮智能浪潮席卷全球的重要历史节点，大会以“智能伙伴 共创未来”为主题，共同探讨如何促进全球人工智能向善普惠发展。

大会认为，人工智能全方位重构人类经济社会发展模式，将对人类文明进步产生深远影响。人工智能带来的机遇和挑战同步上升，一方面，人工智能深度融入产业、民生、科研等诸多领域，为经济转型、技术突破、效率提升注入强劲动能。另一方面，在人工智能持续迭代过程中，技术伦理、数据安全、监管体系等问题也随之凸显。

大会认为，人工智能是全球治理的最新课题，要秉持“全球共济”精神，通过人工智能国际协调为全球治理树立新范式。应当统筹发展和安全、平衡创新和治理、兼顾本国利益和全球福祉，坚持以人为本、智能向善、公平普惠、协同共治原则，加强人工智能国际合作，讨论和解决智能时代人类共同面对的

重大问题、紧迫挑战，把人工智能能力建设国际合作置于人工智能全球治理的突出位置，着力解决资源分布不均、权力分配不公等突出问题，帮助全球南方国家在人工智能发展进程中平等受益，助力落实联合国 2030 年可持续发展议程，确保人工智能造福全人类。

一、创新是人工智能发展的核心驱动力。我们倡导培育以企业为主体、以市场为牵引、以应用为导向、以科研为基础、以人才为根本的创新生态。人在人工智能发展中发挥关键作用，要高度重视培养具备数字和人文素养、创新能力和全球视野的人工智能领域人才。

二、推动人工智能赋能千行百业、走进千家万户。人工智能具有鲜明赋能属性，应当积极推广“人工智能+”模式，激发全球创新活力，探索全新科研范式，打造智能经济新形态，培育以大模型为基础、智能体驱动为引领的词元经济新模式，协同推进传统产业的改造升级、新兴产业的培育壮大和未来产业的谋篇布局，推进人工智能在科学、教育、工业、农业、卫生、交通、气候、环境、政务等领域的应用。

三、开源开放是人工智能普惠发展的重要路径。应当以负责任的方式鼓励共建开源生态，积极打造开放包容的国际开源社区，尊重企业自主选择，在重视知识产权保护的前提下，促进人工智能研究成果和技术经验的分享，提升人工智能技术和服务的可及性。

四、数据在人工智能发展中发挥重要作用。应当把握数据高流动性、高赋能性特征，深化数据资源开发利用，切实保障数据安全，共同加强个人信息保护，加快构建数据产权等基础制度，确保数据可管、可控、可追溯，促进数据安全有序流动和高效开发利用。

五、发展可持续的人工智能，推进人工智能与能源系统的深度融合。通过在新能源富集地区布局规模化智算集群，引导算力负荷与电力资源的协调调度，关注人工智能对能源消耗和生态环境的影响，推动人工智能绿色低碳发展，促进人工智能与能源双向赋能。

六、关注人工智能技术对社会转型的影响，系统评估和应对人工智能对就业结构的冲击，发挥人工智能在创造就业岗位方面的潜能，积极开展人工智能教育和技能培训，激发全社会创新创业活力，切实维护劳动者的权利和尊严，构建更有温度的智能社会。

七、重视人工智能带来的安全风险，积极推动构建法律法规、技术监测、风险预警、应急响应体系，坚持风险导向、敏捷治理，探索分类分级管理，筑牢安全底线，防范滥用恶用，确保人工智能始终处于人类控制之下。

八、人工智能前沿技术的升级迭代给网络安全带来复杂影响，对金融、电力、通信、交通等关键基础设施构成潜在威胁。全球领先的人工智能企业应以审慎态度推进人工智能研发，为

人工智能大模型加装必要护栏，确保前沿模型的安全部署，防范人工智能带来系统性风险。

九、智能体是人工智能产品和服务的新形态，应当明确决策权限和行为边界，建立行为追溯和风险提示机制，着力提升智能体内生安全能力，化解应用衍生风险。

十、应当共同防范和打击恐怖主义、极端势力和跨国有组织犯罪集团对人工智能技术的滥用恶用，在国际上建立必要的危机管控和应急响应机制。

十一、应当践行真正的多边主义，发挥联合国在全球治理中的核心作用，重视人工智能全球治理碎片化问题，支持联合国全球人工智能治理对话、国际人工智能科学小组两项机制发挥作用，进一步凝聚全球共识、分享最佳实践，确保各国权利平等、机遇平等、规则平等。

十二、探索共同构建人工智能国际经贸规则，促进人工智能产业链供应链普惠合作，支持发展中国家更多参与全球分工，确保各国企业公平参与国际竞争，维护人工智能产业链供应链稳定。

十三、针对联合国可持续发展目标执行严重滞后的实际，应当持之以恒推进落实“加强人工智能能力建设国际合作”联大决议，结合发展中国家的实际需求开展务实合作，支持各国发掘自身潜力和独特优势，共享人工智能发展红利。

十四、加快安全、产业、伦理标准制修订，兼顾技术进步、风险防范、社会伦理，建立具有广泛共识、科学、透明、包容的标准规范框架，促进各国间标准互认，降低人工智能国际合作的成本和门槛。

十五、尊重各国历史文化、社会制度、文明形态差异，保护世界文化多样性。推动人工智能技术适配多元文化场景，促进各国文化的创造性转化和创新性发展，维护世界文明多样性和各国文化独特性。

大会期间，《关于成立世界人工智能合作组织的协定》签署仪式在上海举行。世界人工智能合作组织是全球首个人工智能政府间国际组织，总部将设在上海。世界人工智能合作组织坚持向善普惠原则，旨在支持全球南方国家加快人工智能发展，提升创新、应用、治理综合能力，追赶科技革命浪潮，弥合智能鸿沟。世界人工智能合作组织的基本原则是遵循《联合国宪章》宗旨和原则，尊重各国主权和文明多样性，坚持平等相待，践行多边主义，秉持共商共建共享理念，实现互利共赢。世界人工智能合作组织向所有国家开放，并将同其他国际组织保持沟通，探讨适当的合作形式。

来源：世界人工智能大会

<https://worldaic.com.cn/news/a7a32ea4f41f4e86937f51ec201e84b4?tab=0>

（三）AI 发展核心拐点，从静态数据走向经验驱动（2024 图灵奖得主、“强化学习之父”、阿尔伯塔大学计算机科学教授 理查德·萨顿）

演讲核心观点：

作为基础理论领域权威，萨顿围绕智能本质与行业发展瓶颈提出三大核心观点：

1. 行业普遍混淆算力规模与原生智能。当下大模型、多模态生成仅完成人类历史数据的模式拟合，仅能复刻既有知识，不具备自主发现新知的能力，事实偏差、逻辑缺陷等问题难以根除。市场对 AI 存在过度炒作与过度恐慌两种极端视角，全球行业应当坚持多边合作、互利共赢。

2. 呼吁建立统一的通用心智科学。综合心理学、图灵原始理论、工程实践多重定义，他提出智能的核心是依靠动态行为适配环境、完成目标的自适应能力；强化学习是贯通人类、生物、机器智能的核心方法论。大众熟知的图灵测试并非图灵初衷，机器自主交互学习才是其核心构想。

3. 静态标注数据模式已触达天花板，经验驱动是下一代 AI 核心路线。人工标注资源存量持续衰减，无法支撑模型长期迭代；而智能体依托自身与环境交互产生观测、动作、奖励信号自主学习，数据随自身能力同步扩张，AlphaGo、奥数推理系统 AlphaProof 均印证该路径可行性。现有大模型缺少目标

与奖惩反馈，无法自主校验信息真伪；经验型智能体可依靠环境反馈持续校正认知，自主进化是不可逆的技术大势。

来源：世界人工智能大会

<https://mp.weixin.qq.com/s/1NVWR2wvP7T1wLsGZCtHpQ>

演讲全文：

理查德·萨顿：非常感谢，欢迎女士们、先生们，我非常高兴今天能来发言，同时分享一下我对人工智能 AI 未来的想法。讲到未来，现在我们应该怎么样，今后应该怎么样，应该走向什么样的方向，首先我想先说两句，我在前面关于治理的开幕式之后，我觉得我是非常支持合作的观点，共赢的观点和建立合作伙伴关系的观点。

今天我主要和大家谈一下我们怎么思考人工智能，特别是每个人都觉得人工智能现在在发生快速的变化，并且有很快的进步，每个人觉得是一样。但是如果真的是这样你也要思考一下是不是夸大或者夸张，可以看到人工智能背后的推动力可能也是夸大了它的进步或者重要性。

有的人会害怕 AI，他们也有原因和理由夸大重要性，特别是它的进步到底有多快，它到底真的有快速真的进步吗有一些方面、一些事情真的如此。

比如说可以看到现在有很多的突破，我觉得应该是科学上的突破，这些科学上的突破出现在了我们的机器可以非常熟练使用语言这方面的能力，同时可以生成视频、图片。

这些东西也是催生了很多的新行业和产业，也给我们带来很多的经济上的应用，带来真正的经济上的价值，这些非常重要。但是感觉好像这些还是比较大规模的运用还是模式的识别，我们可以把智能和计算两者不能混淆，这里大多数是计算的能力所以我们要特别关注这一点，必须要把这两者的定义区分开。

讲到 AI 的科学就是智能的科学，它真的在快速进步吗，我觉得并不是如此，我们现在对真正体系系统的理解是很少的，待会看所有 AI 模型的时候，这些系统其实它们主要是在使用人类知识的力量，并且把它再交付给我们，他们并没有能力发现自己的知识，从很多方面他们是比较弱的，他们在整个思考过程中还是存在一些问题的，并不是如此强大，这给我们带来一个问题什么是智能我并不指望在这里告诉大家什么是智能，但是会给大家一些观点可以引发大家的思考，同时思考一下你同意哪一种观点和定义。

刚刚开始讲智能的时候到智能是什么意思，就是威廉·詹姆斯(William James)，他是心理学的奠基之父，这是超过 100 多年前的事情了，他对于智能的定义可能并不一定是说的智能而是所谓的心智，他说所谓心智最大的标准是要达成一定的目标，但是使用多样不同的手段。

就我听来这是一个目标，要有相同的结果，但是用不同的手段实现，我们要通过心智，通过智能实现一个目标。

或者可以使用其他的含义或者定义，就是要像人一样行为，通常是通过图灵测试进行表达，但是图灵从来没有用过图灵测试的提法，他讲的是模仿游戏，所以可能大家普遍有一种误解，图灵认为的像人一样的形式不觉得这是人工智能的测试。

不过我们可以看到现在已经普遍接受了，现在有大语言模型在进行训练的时候要让它们像人一样进行行为，这也是其中的一个含义，我们有时候也会使用，但是这和威廉的定义不太一样了，如果我们查字典，一般的字典可以看到它的定义是获得和使用或者运用知识与技能的能力。

它是指获得或者是运用知识与技能的能力，我们如果再回到人工智能之父，他说智能是通过计算来达成目标的能力，这里又提到了目标，又提到了计算，但我们讲的是人工智能，所以甚至就算是自然的智能也是需要计算，这些是各种各样不同的智能定义。

我也有智能的定义，通过行为的适应来实现目标的能力，不断进行适应和调整。除此之外我个人的观点，我是觉得我们应该要设立关于心智的科学，不仅仅是自然的心智或者技术的应用，而是所有综合性的心智科学，他它是可以用于人类、动物、机器所有的这些心智都有共同的特征，都有它共同的特点，同时也都是为了要实现目标，随着时间的推移采取行动，但未来可能更多的心智是机器的心智。

现在没有哪一门现有的科学把这些全部包含进去，无论是心理学、人工智能、认知科学，都没有办法做到这一点，所以我们没有办法把所有的机器都包含进去。因为我个人比较感兴趣的强化学习是综合性心智科学的开始。

我今天讲的主题是我们现在还处于人类数据的时代，所有的这些 AI，它们都是培训预测人们的语言下一个词，人们的标签，并且也是由人类专家进行微调的，大部分的机器学习是要进行知识的转移，从人转到机器身上。

这种方法现在已经达到极限了，很多高质量的数据源被用完了，生成新的知识是这个范式没有办法实现没有办法做到的。我们现在在迈入所谓的经验时代，AI 需要新的 U 数据来源，它能不断增长不断改善，随着智能体越来越强，它不是静态的数据集，因为静态的数据集是不够的。

我们要得到这个数据源它是来自智能体，它自己的经验，它和世界的互动，我说的是它的经验是第一视角的。我觉得我们这些数据的信号它不断来来回回进行传递，这是非常快速，是在智能的智能体和它创造的环境就是世界当中来传达的，这些信号它的经历、行为、观察，这些来定义它的目标到底是什么，所以这是人和动物学习的方式。

心理学研究已经告诉了我们了像 AlphaGo 是通过这样的方式来下棋，也是最近 AlphaProof 系统在国际奥数比赛当中得到

比较好的成绩，接下来给大家看一个视频，也看出一下这是一个怎样的过程。

这是婴幼儿，他在玩自己的玩具，大家可以注意一下他的数据，他收集到来自世界上的这些数据，其实是完全取决于他的这些行为，他做些什么，他会从一个玩具到另一个玩具，从每个玩具身上进行学习，玩完一个之后不玩了，玩下一个。他的行为是决定了他的输入，他是没有静态的数据集，不是事先构架好的，所以这些数据在打造的过程中是最为适合他的需求，以及也适合他的认知能力范围和水平，这是最关键的信息。

然后我们可以看一下无论是动物或者人他进行各种行为的时候，这是我们做的，现在我们都坐着，但是我们的身体是采取行动而且是快速采取行动，对周遭发生的事情反应。

我们是高带宽的信息处理，信息来了足球运动员马上就要决定，所有这些数据来了之后，他接下来应该怎么做，什么是比较明知的踢法。还有打棒球的人他击球的时候，当球来的时候，我们可以看到就这么短的时间内他马上要决定怎么来进行挥杆。还有鸟，可以看到还有狮子，就算是人在这里交谈，是很多的处理，也是为了完成一个目标，同时也是基于我们的经验，希望给大家有所理解。

再看一下这张幻灯片的内容讲到经验型的 AI，智能体是交互这些信号它的行为，它的感知，这些是它的经验，这可以说是我们的焦点，这是智能的全部，智能就是要打造这样的互

动，因此我们可以看到这是一个焦点或者重点，是所有的心智终点，要采取行动要采取行为，并且有所感知。

可以认为一个智能体它是智能的话，它能够预测和控制它的经验，它的感知，如果说它没有经验，那么我觉得根本没有什么智能可言。

特别是我们没有办法去说这种行为的方式比另外一种行为的方式好，我们必须要有经验，要有奖励信号告诉我们这个行为得到奖励，但是另外一种行为没有让我得到奖励。现在的大语言模型它是没有这种奖励的，它没有办法知道这个行为是好的还是不好的。

所以这最终是为什么它们有很大的限制，因为它们没有目标。

观察也是很类似的，这就是所谓的事实准确性，大语言模型基本上没有办法把真的和假的区分开。如果没有经验，你的知识是预测去发生了什么，你没有经验的话，你可以看到之前发生的，但是你没有数据没有经验作为输入来源，你就没有办理知道你做出了预测到底是对的还是不对的，但是有了经验的AI就有奖励、目标，有真实的情况，就是你认为会发生的事情真实发生了。

这是我们的想法，可能对大家来说是显而易见的，如果我们在研究AI的话，大家可能觉得这是非常自然的，这句话是来自 Alan Turing 的，他是计算机科学奠基之父。这是 1947 年

的时候他说的，这时候甚至没有 AI，他看的是自然的智能，他说我们需要的是一台机器，它可以从经验当中来自主学习，这是非常基本但是又是非常深刻的想法，它没有成为关键的现代 AI 组成部分。

但是我觉得这个想法现在慢慢变得越来越重要，特别是在机器人当中，在一些工业的应用当中，所以我也很喜欢这样的说法，是维克多·雨果说的，他说“恰逢其时的思想，势不可挡”。

所以我觉得我们要从经验当中进行学习，这个想法是所谓的恰逢其时，可能这个时代还没有真正到来，这是我们在迈入这样的时代。

做一下总结，我觉得 AI 现在终于转向从经验来进行学习，而不是从人类的数据进行学习，这样它会变得更加强大，因为它可以持续学习新的东西，虽然现在有很多的炒作，甚至是对 AI 的恐惧，现在的 AI 智能还不够强大，我个人认为它是比较弱小而且是不可靠的，因为会产生很多的错觉，它可能告诉我们的东西是错的，但是同时它是非常有用的，它也给我们带来催生了很多的产业，而且所有人可以使用，每个人都可以使用，所以大家都感到非常激动，因为大家都发现 AI 现在还不是真正的智能，或者真正的自主智能，还没有达到这一步，但是我们在迈向这样的时代，我们要习惯，所以现在是第一次和它们接触，还没有来到真正超级智能时代，或者智能增强人类的时

代，但是会给我们带来巨大的变革，应该是由人类推动的，同时这个变革会给到人类。

所以可以看到如果作为一个从经验当中学习的人，我自己本人就是从经验当中学习的，一生都是如此。我觉得我可以预测到经验时代的到来，所以我们要说的是欢迎大家来到经验时代，非常感谢大家的聆听，谢谢。

来源：图灵人工智能

<https://mp.weixin.qq.com/s/UKlbKANlhdAn1XKGpaxo8Q>

（四）当智能体进入物理世界（阶跃星辰董事长、千里科技董事长印奇）

演讲核心观点：

立足产业实操视角，印奇梳理 AI 迭代周期，并预判未来 3-5 年产业结构性变革：

1. AI 产业分三阶段演进，物理场景通用智能体是第三波核心主线。第一阶段为预训练大模型+对话应用，第二阶段为强化学习推理+代码开发辅助，第三波具身智能体带来的产业变革，影响力将超过 PC 互联网与移动互联网变革之和。当前行业已临近 AGI 临界点，代码智能体可独立完成多日完整开发，编程语言成为人机通用交互媒介，智能体是通用 AI 时代最优产品载体。

2. 三大底层变革机遇将重塑数字产业：一是大模型原生跨设备操作系统，统一调度算力、数据与任务；二是全品类硬件终端重构，PC 转为个人算力工作站、手机成为全域设备中控、新能源汽车率先落地机器人形态；三是全新人机共生网络，打通人、智能体、物理终端交互链路，重构价值分配与交易规则，兼容 Web3 等前沿技术体系。

3. 产业设计逻辑转向人机协同共生，AI 是生产力工具而非人力替代者。行业需搭建智能体可信可控约束体系，推进技术普惠；同时新技术将催生全新职业赛道，诞生单人即可完成全业务流程的超级个体。

来源：世界人工智能大会

<https://mp.weixin.qq.com/s/INVWR2wvP7T1wLsGZCtHpQ>

演讲全文：

非常高兴能在 WAIC 的开幕式上，和大家一起探讨人工智能技术和产业发展的下个阶段。过去 5 年，AI 快速发展，回看每一次 AI 爆发，都离不开两股力量的结合：超级模型与超级应用。Transformer 大语言模型和 chatbot 聊天机器人的结合，带来了 ChatGPT 时刻；强化学习、推理模型与辅助编程的结合，打开了 AI 编程的巨大机遇。那么，下一波爆发点会是什么？我们看到了一个巨大的可能，就是智能体模型和下一代终端产品的结合。所以，今天我想分享的主题是：当智能体走进物理世界。15 年前，我刚刚大学毕业，选择的第一份工作就是成为一名 AI 创业者。15 年后，AI 创业从一个小众赛道变成这个时代全球最大的共识之一。我深感能在这个时代成为一名 AI 创业者，无比幸运。当今的每一位 AI 研究者和 AI 创业者，都是站在历史巨人的肩膀上。一边是 70 年来几代 AI 科学家孜孜不倦的探索。不管是符号学派还是连接学派，再到深度学习和大模型。细看整个 AI 技术的发展脉络，都呈现出高度的连续性，同时又具备从量变到质变、如阶跃函数般的跳变性。另一边，沿着产业和应用的角度，IT 行业、PC 互联网、移动互联网、云计算、半导体、新能源，共同构成了 AI 爆发不可或缺的序章。2026 年，模型能力跨越了能力的临界点。可以说，我们已

经站在 AGI 高峰的山脚下。5 年前，AI 能连续执行的任务时长仅仅是几秒钟，今天它已经能够独立工作数十小时，而且进步的斜率还在不断增大。我发现其中一个有趣的规律，AI 发展的里程碑似乎总和一个词紧密相关，那就是“语言”。从早期的图灵测试，机器能否如人类一般掌握自然语言，是衡量智能最重要的标准；后来大语言模型 `predict next token` 的模式出现，一开始大家认为这不是智能，到后面惊讶的发现这恰恰开启了通往智能的真正大门。再到 2026 年的关口，AI 已经不光掌握了人类交流的语言，还在人类最重要的智力劳动——编程上快速追赶甚至超越人类。这可能会成为图灵测试的 2.0 版本：在最重要的语言框架下，机器是否能超越人？下一阶段更关键的问题也将变成，我们如何为 AI 进入物理世界设计一套新的语言体系？哲学家维特根斯坦曾说：“语言的边界，就是世界的边界。”在 AI 的世界里，也似乎遵循这个规律。无论是从模型技术演进还是产业发展的角度，“智能体”都变成 AI 下一阶段最重要的关键词。回看历史，所有的产品都是某种服务的载体。在信息革命 80 年的时间里，最重要的产品形态，从大型计算机、个人计算机，再到软件、网站、移动应用，接下来可能快速转变成“智能体”这样一种新载体。智能体，将不仅是一个聊天的机器人，它能感知外部物理世界，理解用户深度记忆和偏好，综合决策和执行，完成用户交予的各项任务。智能体，正在成为生产力的最小单元。未来，每个人身边都会有专业智能

体：工程师有代码智能体，设计师有创意智能体，研究员有分析智能体……这些智能体，会让一个人拥有一支团队的能力，也让屏幕里的智能真正走进我们的工作和生活。因此，智能体带来的不是一个单点应用机会，而是三个结构性变化：新系统、新载体、新网络。第一，是新系统。过去四十年，每一次计算平台的变化，都伴随着操作系统的变化。PC 时代管理软件，移动时代管理应用，智能体时代则要管理很多能够理解目标、规划路径、调用工具并执行任务的智能体。因此，智能体能做什么，不只取决于模型有多聪明，也取决于操作系统为它打开了多大的能力边界。因此我提出一个公式：智能体能力=模型能力×智能体操作系统能力。如果说模型决定智能体能“想多远”，那么 **Agentic OS** 决定智能体“走多远”。它连接模型与数据、工具、接口和设备，让智能体不仅能够理解世界，也能够采取行动。这就是智能体原生操作系统，它将成为未来人机交互和产业生态的重要基础设施。第二，是新载体。移动互联网时代，终端连接人与数字世界；智能体时代，终端将带着智能走进物理世界。它不再只是等待人的操作，而是能够感知环境、理解情境并采取行动。这意味着，终端的设计逻辑正在改变：过去，我们围绕人的操作方式设计机器；未来，我们将围绕人与智能的协作方式重新定义终端。终端的设计理念也将从“以人为本”，进一步走向“人机共生”。电脑将成为 **Agent** 工作站，手机将成为随身智能中枢，汽车将成为移动智能空间，机器人

将成为人类身边的智能伙伴。更重要的是，这些终端不再彼此孤立。电脑、手机、汽车和机器人，将成为同一个智能体在不同场景下的不同身体。智能可以在多个终端之间连续存在、自由迁移、协同完成任务。未来的终端竞争，不再只是设备之争，而是谁能让智能体更好地感知环境、采取行动，真正进入物理世界。第三，是新网络。互联网时代，连接的是人、信息和服务；智能体时代，我们将还将连接能够理解目标、调度资源、协同工作的智能体。这就是 A2A 网络。从产业视角看，这不仅一种通信协议，更是一张由 Agent 和人类共同构成的人机共生价值网络。未来，数十亿人、数百亿设备和智能体，都可能成为其中的节点，形成一个规模更大、协作效率更高的新型网络。对创业者而言，这意味着新的产业机会：智能体会有自己的身份和信用。人们可以找到它、了解它擅长什么、把任务交给它，也可以根据它的表现付费、评价，以及与它合作。当智能体能够自主寻找伙伴、组织协作、完成交易，A2A 就不再只是技术协议，而将成为智能体经济的基础设施。但也正因为智能体正在进入组织、网络、设备和真实世界，我们必须看到：人机共生带来的，不只是能力跃迁，更是秩序重构。当智能体开始代表人行动，我们必须回答三个问题：它可以代表谁行动？行动产生后果，谁来负责？当它连接数据、设备和真实环境，怎样确保身份可信、权限可控、行为可追溯？过去，我们更多是在科幻作品中思考这些命题。今天，随着智能体走进现

实，它们正成为必须真正面对的问题。为此，阶跃正与图灵奖得主姚期智先生领衔的上海期智研究院开展前沿研究。我们也期待与全球专家学者一道，探索智能体时代新的技术体系、治理机制与社会契约。作为年轻一代 AI 创业者，我们既是这场技术浪潮的参与者和受益者，也是下一阶段秩序的共同建设者。我们有机会把智能从屏幕带入真实世界，也有责任让它更可信、更安全、更普惠，成为人类长期的伙伴。Agent 时代最大的意义，不是替代人，而是放大人。它正在孕育新一代 AI Native 建设者，也催生新一代 AI Native 组织。AI Native，不只是会用 AI，而是从第一天起，就把 AI 写进产品、系统和组织的基因里。对个人，它让一个人拥有一支团队的能力；对创业者，它意味着创造新产品、新组织，甚至一个全新的产业。每一次技术革命，都会催生新一代伟大的企业。我们相信，人工智能时代的超级建设者，也一定会从 AI Native 组织中诞生，并由他们参与书写下一段产业历史。最后，我想回到阶跃最朴素的使命：智能阶跃，十倍每个人的可能。我相信，未来不是机器替代人的世界，而是人与智能共同进化的世界。在 AI 行业的 15 年里，我觉得最有意思的地方是：技术还年轻，产业还年轻，很多问题还没有标准答案。大家不是在一条已经铺好的路上竞速，而是在一起把这条路修出来。因此，我们这一代 AI 创业者要做的不只是创造更强的模型，更要让智能走进真实世界，

成为每个人都能使用、都能信任的力量。当智能体进入物理世界，每个人的可能，都将被十倍放大。

来源：人工智能 X 机器人

<https://mp.weixin.qq.com/s/JAtjWYSG05UbuFthzuKo7A>

兴信咨询设计研究院整理

(五)要想驾驭 AI 和信息革命，战略制定要具备韧性(图灵奖获得者、美国康奈尔大学荣休教授、中国科学院外籍院士约翰·爱德华·霍普克罗夫特)

演讲核心观点：

人工智能推动信息革命，影响将会远超工业革命。他提到，2012 年前，人工智能还只是一个小众的学术分支；大约在 14 年前，一项图像识别大赛中的成果，将人工智能从小众研究推向了创造数百万就业机会的主流学科；到了 2015 年，研究人员开始致力于探索“压缩互联网内容”，这一尝试则推动了大语言模型的爆发。

“人工智能很有可能还将持续发生重大变化。各国需要制定相应的战略来驾驭这场革命，而相应的战略必须要具备适应变化的韧性。”在霍普克罗夫特看来，教育能帮助人构建这种韧性，并培养新型人才。他说：“过去，黄金、石油和农业铸就了国家的强盛；而在未来，人才是一个国家的核心竞争力，只有那些真正致力于人才培养的国家，才能在信息时代占据领先。”

来源：中国青年网

https://news.youth.cn/jsxw/202607/t20260718_16771783.htm

(六)物理智能的使命是“把人还给人”“可靠运行”是行业重心(复旦大学浩清特聘教授、通用物理智能研究院院长苏昊)

演讲核心观点：大模型要走出幻觉，必须先走出“数字世

界”的边界，去亲身体验，把自己交给“物理世界”去裁判——做出预测，采取行动，被现实修正。这个古老的过程叫作实验，也是物理智能的关键。

物理智能真正缺的，是一个能聚合人类全部物理世界知识的模型。从认知的源头数起，物理知识至少有六层阶梯，层层筑基。

下三层，属于客观世界。一是关于物体的知识，二是关于状态的知识，三是关于动力学的知识。上三层，则是下三层绑定了“主体”之后的样子。四是关于功能的知识，五是关于目标的知识，六是关于行为的知识。

来源：腾讯网

<https://news.qq.com/rain/a/20260718A07BA500>

演讲全文：

尊敬的各位领导、各位嘉宾，大家好。很荣幸在世界人工智能大会分享一些个人观点。本论坛的主题是“基石”，而我今天想聊的，恰恰是智能最底层的那块基石——物理智能。

先说结论：对物理智能，我是坚定的乐观派。看清幻觉，是为了让现实来得更快——这就是今天的标题：“从幻觉，到现实”。

一、从大模型到物理智能

大模型的进步有目共睹，但为什么最聪明的模型也难免幻觉？我认为最本质的一条是：语言只是世界的投影。人类先在

物理世界摸爬滚打，才把经验压缩成语言；模型学的一直是这道影子，从未见过投下影子的实体。它能流利地描写“杯子掉在地上会碎”，却从没有机会感受一只杯子的重量。知识没有现实的锚点，说错了，它也无从自知。你可以在屏幕里雄辩，但你说服不了重力。

大模型的幻觉，很大程度上，就出在没有“身体”。要走出幻觉，必须走出“数字世界”的边界，去亲身体验，把自己交给“物理世界”去裁判——做出预测，采取行动，被现实修正。这个古老的过程叫作实验，也是物理智能的关键。

二、物理智能真正缺什么：把散落的知识聚合

那物理智能真正缺的是什么？

我们需要的，是一个能聚合人类全部物理世界知识的模型。从认知的源头数起，物理知识至少有六层阶梯，层层筑基。

下三层，属于客观世界——有没有“我”，它都在那里。

一是关于物体的知识：世界，是由一个个独立、持存的物体组成的——一只球滚到沙发底下，看不见了，但它还在。先认清“有什么”，其余才谈得上。

二是关于状态的知识：这些东西，此刻是什么样子？球在沙发底下，它有多重、是软是硬，都在这一层。

三是关于动力学的知识：世界自己会怎么变？球会滚，水会流，松手的东西会往下掉。而“力”就在这一层——你用眼睛，看不见它。

上三层，则因“我”而生——每一层，都是下三层绑定了“主体”之后的样子。

四是关于功能的知识：这东西，对我有什么用？把手是用来握的，杯子是用来盛的。同一把椅子，对人“可坐”，对蚂蚁不是。物，因为主体，才成了“器”。

五是关于目标的知识：我要世界变成什么样？球在沙发底下，我要它回到手里。球在哪，是状态；我想要它在哪，是目标。目标状态的“是”与“应是”，一字之差，隔着整个主体。

六是关于行为的知识：知道要做什么，和手真能做到，是两回事——端一杯水穿过房间，不洒，是分寸；系鞋带、用筷子，是巧劲。这些本事，不写在书上，长在手上。

从物体到行为，越往上，越不是“看”会的，越得在物理世界中亲手去“做”。每个婴儿的头两年，就是沿着这道阶梯一层层爬上来的——皮亚杰称之为“感知运动期”：人类智能的地基，是用手打的。

但模型没有这样的童年——它学东西，主要靠人类留下的记录。麻烦在于，这六层知识散落在互不相通的载体里，越往上，被记录下来的越少。互联网视频最海量，却基本停在阶梯下段——样子与运动看得见，力与手感够不到；教科书方程最精确，写的正是动力学，却只针对理想化的世界；真机数据有力觉、有操作示范，直抵上段，却少得可怜。

语言模型是幸运的：互联网早已替它把语言知识聚合完毕。

物理知识没有这份幸运，波兰尼说，我们知道的，远多于我们能说出的。所以，阶梯的上半截，至今没有被成体系地记录下来——光靠刷网页，刷不出完整的物理智能。聚合，就是我们这一代人要完成的科学工程：把视频的广度、方程的精确、真机的真实、直觉的细腻，熔进同一个模型，互相校准补盲，把阶梯补完整。物理智能要从幻觉走到现实，中间隔着的，正是这道阶梯。

三、时代的刚需，和它到来的方式

这道阶梯，为什么非爬不可？

答案不在语料里，不在机房里，而在现实世界。

今天的 AI 会写诗、会写代码、会做 PPT——但它帮不了一位老人翻身。智能的价值大多还停在数字“比特”世界，可人最沉重的那些需求，全在物理“原子”世界。再看身边：老龄化写在日历上。需要照护的人日益增多，能投入照护的人手却在减少。另一边，高空、井下、高温环境中危险、繁重的作业，也面临劳动力短缺的困境。需求都在，缺的是人手。物理智能不是人的替代者，而是人的协作者——填补的是人手的缺口，把翻身、搬运这样的体力活交给机器，把照护者的时间，还给陪伴与关怀；把危险的作业交给机器，让人退到安全线之后，承担判断与创造。物理智能的使命，是把人还给人。

那它会怎么到来？作为技术人员，说说我的个人判断。它不会在某场发布会上一夜“实现”，更像当年的电气化——先点

亮工厂和仓库，再走进商店和医院，最后才走进千家万户。物理智能的价值，不必等到终点，沿途就在释放。

这条路上最险的一段，是演示与产品之间那道可靠性鸿沟。填平它，靠的不是热闹，是基础工作，是积累数据——尤其是带力、带交互的那一半。填平它，靠的是行业标准，供应链，还有最难攒的社会信任。信任只能靠可靠性一点一点去挣；而这份信任，也是物理智能落地最底层的那块基石。

四、三个判断

最后，我留下三个判断，供未来检验。

第一，物理智能的突破口，不在模型的架构，而在知识的聚合。聚合这件事，注定超出任何一家机构的边界——视频在互联网上，方程在教科书里，力觉数据在各家实验室里，操作的直觉在亿万劳动者身上。没有谁能独自集齐这道阶梯，它需要整个行业乃至全社会的通力协作：共建数据，共立标准，共享仿真、评测等基础设施。当这些知识真正熔进同一个模型，物理世界就会迎来属于自己的“互联网时刻”——所谓 GPT 时刻，只是它的副产品。

第二，行业的重心，将从“演示做得多惊艳”，转向“运行得有多可靠”。工程上有个说法叫“几个九”：从 99%到 99.9%，每多一个“9”，难度都指数级上升；演示与产品的差距，就卡在最后那几个“9”里。通用性是终点，而可靠性才是起点——愿意在“9”上下笨功夫的团队，会走得最远。

第三，物理智能将把 AI 从科学的“读者”，变成知识的“创造者”。今天的 AI 读遍了人类几乎所有的论文，却几乎没有亲手做过一次实验；而新知识，恰恰诞生于实验。当它拥有一双能感知、操作、验证真实世界的手，它就能自己提出假设、亲手实验、昼夜不休地修正。新材料、新药的发现，可能因此快上几个数量级。

从幻觉到现实，没有捷径，靠的不是更响亮的叙事，而是对物理世界的敬畏，和沿着那道阶梯一层层往上爬的笨功夫。物理世界是智能最古老的老师，最诚实的考官，也是智能最终的基石。我们选择把答卷交给它。

谢谢大家。

来源：搜狐

https://www.sohu.com/a/1051557105_121118944

（七）科学发现的革命：让科学数据成为基础模型的原住民（中国工程院院士、之江实验室主任、阿里云创始人王坚）

演讲核心观点：

随着科学基础模型的发展，人工智能正从一个单纯的计算机科学分支，演变为支撑所有科学研究的重要基础设施。”王坚说，人工智能将变得跟数学一样基础了，它不再是一个独立的垂直技术领域。“未来 50 年里，人工智能一定会像数学在人类科学历史上所起的作用一样，对我们科学研究起作用。

来源：新民网

演讲原文：

各位嘉宾，大家好。

我今天想从一个例子讲起。

2024 年 11 月，有一篇文章发表在一本天文学领域的顶级期刊上。它只有一个作者。今天的人工智能论文，动辄有几十个、上百个作者，只有一个作者的文章已经很少见了。更有意思的是，这位作者发现了将近 150 万个此前没有被编目的太空物体。

如果只是这样，大家可能还不会特别吃惊。但如果我告诉你，这位唯一作者是一名 18 岁的中学生，而且他使用的数据来自一颗已经退役的卫星；那颗卫星当初收集这些数据，本来并不是为了完成他后来做出的发现，而是用于监测小行星等其他目标。这个故事就很值得我们重新思考：我们对数据的理解，可能还远远没有到位；同时，科学发现的专业逻辑，也可能正在发生巨大的变化。

他不是一位做了很多年研究、读完博士又继续做博士后的科学家。他只是一个高中生，尽管他是在 Caltech 做了实习工作。这让我想起一本关于科学革命结构的书。范式变革发生的时候，常常会出现一种现象：人们能够从旧数据里发现新问题。这几乎是科学变革中的一个基本逻辑。

大家都知道，伽利略当年用的数据，也并不都是他自己重新找来的新数据。所以，在今天这个时代，我们必须对“数据”有一个重新认识。

一、基础模型为什么重要

在人工智能领域，大约在 2021 年，斯坦福的一组教授写过一篇报告，第一次明确提出了今天大家都在使用的一个概念：Foundation Model，也就是基础模型。

我们今天讲模型、讲大语言模型，但“基础模型”这个概念，是到 2021 年才被这样明确提出的。它的重要性在于，它让我们意识到：数据不是简单地喂给模型的材料，而是模型能力形成的根基。

这也是我今天为什么要讲，科学数据应该被放进基础模型。

今天我们说“大模型”，在中文语境里，很多时候会把“大语言模型”里的“语言”两个字省掉。但事实上，完整地说，大语言模型是建立在文本数据上的基础模型。它的根本能力，来自人类对文本数据的长期积累和处理。

但处理文本这件事情，并不是人工智能出现以后才有的。

我们作为在中文环境里读书的人都知道，中国古代的典籍，最初是没有标点符号的，甚至也没有今天意义上的空格。中文是到一百多年以前，才逐渐开始系统使用标点符号。古拉丁文在早期也是如此，没有空格，也没有标点符号。

没有标点符号会带来什么影响？一个人如果不把这句话

读出来，就很难完成阅读和理解。所以，对文本进行处理这件事，早在人类文明史中就已经存在。用今天的观点看，在文本之间加入空格和标点，某种程度上也可以理解为一种早期的 Tokenization。

今天的大语言模型，把人类对文本的处理带到了一个前所未有的高度。Tokenization 是其中非常关键的一环。

二、科学不应该只停留在论文文本里

再回到科学。

很多年前，我就开始和一些地球科学家讨论：我们怎样用人工智能帮助地球科学？

事实上，早在十几年前，也就是人工智能还没有像今天这样流行的时候，科学家们就已经提出过三个非常朴素、但极其重要的要求。

第一，科学家需要获取所有相关的科学数据。

第二，科学家需要获取相关领域的研究文章和论文。

第三，科学家需要一个基础设施。如果换一个说法，也可以把它叫作公共设施。

今天回过头来看，人工智能其实远远没有满足第一个条件。我们今天大多数所谓“AI 帮助科学研究”的工作，仍然停留在文本层面。也就是说，我们更多是在让模型读论文、读文章，最多再看一看论文附带的代码，而不是让模型真正触碰科学数据本身。

代码是一个很好的例子。

如果你学习编程，最早接触代码的时候，可能也是在一个文本编辑器里。但为什么 **Coding** 在今天变得这么重要？因为在模型里，代码没有被简单地当成另一种普通文本。

大家可以想一想：中文和法语之间的差别，中文和英语之间的差别，可能都没有英语和代码之间的差别那么大。代码作为一种数据，恰恰展示了数据如何驱动 AI 的思考。

所以今天我们要问：当人工智能帮助科学研究时，它到底是在理解科学，还是只是在理解科学论文？

如果只是把科学论文的文本读一读，把论文里的代码看一看，那它仍然没有碰到科学数据本身。它理解的仍然只是人类已经写出来的科学知识，而且是被论文形式表达出来的那一部分。

你可以想象，这后面还有多大的空间。

在地球科学里，70% 以上的重要信息并不在文本中。比如光谱、地震波、地层、化石、岩石样本，它们都不是传统文本。我们常说多模态模型，“一张图胜过千言万语”。但在科学里，可能一段光谱、一段地震波，就胜过千万张图。

这件事非常有意思。

三、什么是科学基础模型

所以，几年以前，我们定义了一个概念：科学基础模型。

所谓科学基础模型，就是建立在科学数据基础上的模型，

而不是建立在科学论文文本基础上的模型。

它的目标非常清楚：让科学数据成为大模型的源头。

这里面当然有一个很直接的工作，就是要把科学数据做 Tokenization。可是最终，我们要做的不只是把科学数据切成 token。真正的科学基础模型，是要把文本、代码以及科学数据放到同一个空间里，放进同一个模型里，让模型能够统一处理这些信息。

从这个意义上说，科学基础模型不是一个 vertical model，不是一个简单的垂直模型。今天我们一讲到某个具体领域，就很自然地想到：这是不是一个垂直模型？但我今天讲的不是这样一个独立的、狭窄的技术模型。

科学基础模型应当改变我们对科学数据的整体理解。

四、让石头“说话”：地球科学里的例子

我很快举一个例子。这个故事还在进行当中。

大家都知道，地球科学很大一部分是在研究石头。石头是不会说话的，但所有这些工作，其实都希望石头能够“说话”。石头说话的目标也很简单：我们想知道这个世界是怎么演进的。

地球已经有几十亿年的历史。生命什么时候出现，恐龙什么时候出现，人类什么时候出现，这些问题都放在一个非常巨大的时间尺度里讨论。这里的时间精度，往往是以百万年为单位来计量的。

和地球科学家打交道以后，我逐渐明白了一件事情：人类

最基本的目标之一，是能够在地球上生活得更长久一点。但要做到这一点，我们必须更好地理解地球的长时间尺度。

这很难。

所以，地球科学家和人工智能合作的一个基本目标，就是把对地质时间的认识精度，从过去常见的百万年尺度，推进到万年尺度。这是几个数量级的提升。大家都知道，在科学里，精度提升几个数量级，是非常重大的事情。

我们和南京大学、中科院南京地质古生物研究所等团队合作，做了一个地球科学模型。这个模型规模并不大，但已经做出了一些很有意思的工作。

要理解这项工作，首先要理解化石。

化石里既有时间信息，也有空间信息。但化石记录有一个重要问题，古生物学里叫 **Signor-Lipps effect**。简单来说，最后一个消失的物种，不一定会留下最后一个化石。一个生物能不能变成化石，本身是随机的。我们找到的某个物种最后一件化石记录，并不意味着那就是它真正消失的时间。

所以，当你根据化石推断年代时，并不能简单地说：它就是在那个时间点消失的。这种效应使得我们对时间的确定变得非常困难。

这恰恰是人工智能可以帮助的地方。

今年 7 月 1 日，在一次地质学会议上，相关团队展示了一项工作：参考十万种不同生物的种类、近两万多个地层剖面

的数据，用模型把这些信息真正放到一个统一的时间轴上。这大概是目前世界上最完整、也最精确的相关时间轴之一。

大家可以想象，如果我们最终能够把百万年级别的精度推进到万年级别，我们对这个世界的理解会变得完全不一样。

我们也很高兴，这项工作近期入选了联合国“科学促进可持续发展国际十年”相关认可清单。大家可以由此看到，人工智能对基础科学的推动，已经不是一个抽象的口号。

五、科学基础模型也需要全球治理

当然，要完成这样一个全球性的事情，治理结构非常重要。

尽管这只是地球科学这样一个领域里的人工智能模型，但我们在将近四年多以前，就建立了一个全球治理结构的委员会。这些专家和机制，对这项工作在全球范围内产生影响，发挥了非常重要的作用。

这也说明，科学基础模型不是某一个机构、某一个实验室、某一个国家可以单独完成的事情。它需要科学共同体、数据共同体、模型共同体和治理共同体一起参与。

六、人工智能会像数学一样基础

最后我想说一句话。

因为有了科学基础模型，人工智能到了一个非常不一样的转折点。

这个转折点就是：人工智能会变得像数学一样基础。它不再只是某一个领域里的独立工具。

大家可以想一想，数学在人类科学史上起了什么作用。

未来五十年，人工智能对科学研究所起的作用，可能也会像数学一样深远。

谢谢大家。

来源：智慧城市行业分析

https://www.smartcity.team/professional/wangjian_waic2026/

（八）无界创新与有界守护——人工智能发展的普惠与安全（商汤科技董事长兼首席执行官徐立）

演讲核心观点：

他提出无界的人工智能创新更需要安全边界的守护，其中既包含产品理念安全、模型发展路径安全，也包括物理边界安全。演讲中，徐立还展示了当前相关产品实践，展现出坚守发展与安全并重的原则，已经成为业内的普遍共识。

来源：世界人工智能大会

https://mp.weixin.qq.com/s/STvNsIfNKytijfJ_7Hl_g

演讲全文：

各位嘉宾，大家好。

今天我想分享我们对于当前人工智能发展，以及如何做好、用好人工智能产品的一些思考。人工智能的发展是无界的，我们看到它正朝着 AGI、ASI 的方向演进，充满无限可能。但如何用好人工智能，核心在于如何为它构建合理的使用边界。

在从专业工具走向大众工具的整套演变过程中，人工智能的使用方式正在发生悄然变化，过去仅面向专业程序员的能力，如今普通大众也能通过多种方式获取原本只有高级软件工程师才具备的技能。

每一次新的生产力突破，往往都会重新定义大量新的工作，尽管现有部分流程会被自动化，但由此拓展出的边界、带来的可能性要大得多。因此我们在探讨如何使用 AI 时，更多关注的是如何拓展 AI 的边界。

在这个过程中，大家通常关注的是 AI 消耗了多少算力、消耗了多少 Token。而我们认为，当 AI 的服务边界真正拓展到超级个体时，未来的核心计费维度应当是任务完成的价格。

因此，衡量单位从 Token 转向 Task，是一种全新的经济学的计算。这部分 Task 的计算成本会持续降低，从而是使得说真正意义上 AI 变成一个普惠的门槛。

很多人会担心部分岗位会消失。我认为确实如此，一些偏流程化、两端都是系统的岗位，一定会被 AI 所优化。但与此同时，AI 拓展的边界会催生大量全新的职业、公司形态与组织形态。

世界经济论坛有一组数据：到 2030 年可能会减少 9200 万个岗位，但同时会创造 1.7 亿个新岗位。

在我看来，这只是很小的一个切面，AI的发展速度远不止于此；前面提到的超级个体，将会演化出新的职业定义，进而赋能整个行业的下游发展。

那么该如何划定这个边界？我们的答案是，必须让AI在安全边界的守护之下发展。

那这安全的理念分成产品理念：不做替代，而是增强；不做依赖，而是让个人成长。

同时要保证模型发展路径的安全，我们在真正预训练、后训练的时候要保证模型训练数据的安全，同时在使用的过程当中，其实也要解决多模型的交叉验证，使得在模型的应用当中的安全。

这是在使用上面的一种范式，那么当然就是物理边界的安全，包括说用沙盒，甚至是具身限定它的使用场景，从而逐步的推动AI在落地当中的这样一个安全边界。

商汤也与联合国相关组织共同推动了「AI Governance for Humanity Lab」（人类AI治理实验室）的建设，也参与联合国全球治理专题报告的撰写，同时与多个国家一道推动基于道德的治理，甚至是教育、AI培训等领域开展了广泛合作。

我们认为，只有打破治理孤岛，与更多合作伙伴一同重新定义下一阶段AI的使用范式，人工智能的发展才能更加健康。

最后，我想用《道德经》中的一句话作结：有之以为利，无之以为用。

今天大家关注更多的是「有」的层面，是模型的发展、参数的迭代、应用的突破；而背后的治理体系、规则框架，是「无」的层面。只有将「有」与「无」并行，那我们才能够做到好好的用。

那做到“能用”，就是让每个人都能用到；“好用”，让 AI 用好的方式被用到；而“敢用”就使得 AI 能够解决安全性的问题。

谢谢大家。

来源：商汤君 商汤科技 SenseTime

<https://mp.weixin.qq.com/s/tdsGCkcSrtFGemOvFmpb4A>

（九）人机共生：AI 向美的实践和未来图景（欧莱雅集团全球首席人工智能官安迪雅）

演讲核心观点：

人工智能并非又一次单纯的技术更新，而是一场前所未有的范式革命。欧莱雅正引航一场美妆变革 – 以人工智能赋能科技、科学和创意技术来释放创新活力，重新定义美妆体验。作为全球美妆行业先行者，我们以雄厚的资源投入为这一愿景提供坚实支撑，投资于自身的未来，以及科技与人工智能的转型进程。

来源：中国新闻网

<https://www.sh.chinanews.com.cn/swzx/2026-07-20/147808.shtml>

（十）AI 赋能破解生命密码——利用前沿人工智能，开发抗疾病与衰老双效疗法（英矽智能创始人兼首席执行官亚历克斯·扎沃隆科夫）

演讲核心观点：生成式 AI 能够高效挖掘基因与蛋白数据，精准锁定“抗病+抗衰”的双效靶点，并完成分子的从头设计与优化。这一创新不仅大幅缩短了新药研发周期，更推动人类医学从“治疗单一疾病”向“全面延长健康寿命”迈出关键一步。

来源：英矽智能

<https://mp.weixin.qq.com/s/o8RCoS yH7zWPC87647KGGA>

二、思想者论坛

（一）AI 的第一性原理：人工智能的威力及其局限（图灵奖得主、中国科学院院士、清华大学交叉信息研究院院长及人工智能学院院长姚期智）

演讲核心观点：

从理论角度看”，为整场论坛围绕 AI 本质的讨论树立了标杆。人工智能的威力无需赘言，它的基石是数学、极限是物理世界，厘清 AI 与它们的关系，才能“立足计算理论与量子物理的本源，探索 AI 科学发展前景，突破 AI 当今极限”。他认为，人工智能今天正推动科学发展，下一个层次是未来由科学推动人工智能抵达新的高度。姚期智教授还前瞻了量子科技的未来，“这是真正的前沿，5 到 10 年内我们将看到量子人工智能的巨大进步”。

来源：世界人工智能大会

https://mp.weixin.qq.com/s/P-RIo8_v0TGcYGQIgqX3-w

演讲原文：

今天我想和大家聊一聊人工智能。我想现在地球上所有人都已经意识到，AI 正在改变世界。它看起来似乎无所不能，所以我们今天要探讨的问题是：AI 的能力是否存在边界？更进一步，是否存在另一个层次的科学认知，是我们还有机会去探索的？我想从理论的视角出发，和大家分享一些思考。

本次演讲的结构大致如下：首先，我会简要回顾 AI 令人惊叹的实践能力，反思 AI 的本质，在此基础上直接回应“是否存在 AI 无法解决的问题”这个疑问，我会给大家展示一些 AI 绝对解决不了的难题。当然，正如我们后面会看到的，AI 存在局限其实是件好事。

其次，我会谈谈我认为未来几年 AI 最令人兴奋的发展方向，这个方向的影响力甚至远超智能机器人带来的改变，将深刻影响社会。我之所以觉得它令人兴奋，是因为它能指引我们找到下一个层次的 AI。要做到这一点，我们不能只盯着“AI for Science（AI 赋能科学）”，更要关注“Science for AI（科学赋能 AI）”，正是科学会把我们带到下一个台阶，而这本身就是一场刚刚拉开序幕的伟大征程。

最后，我会谈谈我认为 AI 研究中值得关注的几个方向。

01

AI 的本质与不可逾越的边界

首先，我们先反思一下 AI 的本质。我们从 AI 高度依赖“机器学习”这个强大工具说起。机器学习作为一种算法范式已经存在很多年，但很长一段时期被主流计算机科学界忽视了，它的解决方案质量一直在持续提升。

机器学习的基本框架是这样的：给定一个任务和模型模板，模板由参数 θ （希腊字母表中的第 8 个字母 Theta）决定，也就是说这是一类算法的集合，我们一开始并不知道该选哪一个。

首先要解决的是表征问题，针对我们要处理的任务，这个模板里是否存在一个合适的 θ ，能让算法完成任务？如果答案是肯定的，或者说我们相信答案是肯定的，那么接下来的“学习问题”就是，能不能在合理的时空成本内找到这个 θ ？和几十年前艾伦·图灵确立的主流计算机科学算法思路相比，这个框架有什么不同？

在标准计算机科学中，这两个问题通常是合二为一的，研究者不需要外部数据辅助，仅通过数学分析就能完成推导。而机器学习最大的威力，也是我们今天看到它取得如此大进展的核心原因，就是把“从数据中学习”加入了算法的武器库。于是AI机器能在围棋上达到冠军水平、能预测蛋白质折叠，这类成果我们已经见怪不怪了。

现在很多人担心AI的能力过强，甚至担心它会失控，所以问一句“AI的能力有没有边界”是非常合理的。除了学术层面的兴趣，搞清楚这个问题对我们设计安全的AI系统也很有帮助。如果我们知道AI的局限在哪里，反而可以利用这些局限，去寻找更好的算法和更安全的方案。

如果你是计算机科学家，可能立刻就能给出答案，甚至能拿出证明：那个著名的“停机问题”依然是无法解决的。关于“停机问题”我就不做技术解释了，本质上就是说我们很难判断一个程序是否会正确运行。为什么这个问题依然超出AI的能力范围？因为我们之前说过，AI本质上就是加装了一组数据的

图灵机，哪怕增加了数据，它依然是图灵机，所以不可能解决“停机问题”。

不过更有意思的是，有哪些实际问题是 AI 无法解决的？哪怕 AI 能下围棋、能预测蛋白质结构，看起来几乎无所不能，你可能会觉得所有实际需求 AI 都能搞定，但事实并非如此。

我举两个例子：第一个和数据安全有关。我们都知道，最新的强 AI 模型已经可以被用来发起网络攻击，你可能会觉得这类 AI 能破解所有安全密码。这对我们来说会是灾难性的，毕竟我们靠加密系统保护银行账户、隐私等等。但第一个例子就是：存在一些加密方法，能抵御所有攻击者，不管是标准算法还是 AI 系统。

我来具体解释一下这个问题，我们看一种特定的攻击模型，叫“选择明文攻击”。假设我们有一个加密系统：我希望朋友能给我发邮件，就算信息在传输过程中被人截获，对方也不知道内容是什么。这个系统的逻辑是：我有一个解密盒，本质是一段程序；朋友有我给的加密盒，能把明文转换成密文，密文在公开网络上传输，任何人都能看到，但只有我能用解密盒还原出原始内容。这个逻辑不难理解，但假设攻击者拿到了你的加密盒，能任意选择明文让加密盒加密，比如他今天输“今天下雨”，得到对应的密文；明天输“我有一只猫”，又得到对应的密文，反复做 1 万次，拿到 1 万组“明文—密文”对。问题由此

产生：攻击者有没有可能算出我的加密盒是怎么构造的？换句话说，在实践中，攻击者能不能算出我的解密密钥？

这个问题非常实用，它对应着我们广泛使用的公钥加密系统：公钥的所有者把加密密钥放在邮箱、网站甚至公开目录里，任何人都能拿到它给我发加密消息。这种“选择明文攻击下的安全性”是公钥系统等核心基础设施的关键指标，像 ElGamal 密码这类加密体系，在特定假设下，从数学上是能抵御选择明文攻击的。

密码学是理论计算机科学中非常深厚丰富的学科，属于计算机科学的数学研究分支。我简单铺垫一下：理论计算机科学从 20 世纪 60 年代起步，之后几十年自然生长，最棒的地方在于新思想不断涌现，而且往往来自外部数学家、生物学家、电气工程师都会带来新思路，每隔 10 年就有让人兴奋的新理论出现。随着计算技术的进步，比如图形处理、互联网、AI、量子计算等等，它和其他学科（物理、生物、经济学）交叉融合，催生了新的交叉领域，也产生了重要的实际应用。

我再举第二个例子，也是解决安全通信问题的，叫“量子密钥分发”。这种方法能让两个从未谋面的人，只通过电话沟通，最终得到一个只有他们俩知道的随机密码，哪怕攻击者全程监听了所有对话，也猜不到这个随机值是什么。

这听起来简直像奇迹，但确实存在可行的方法，而且这个方法超越了传统计算机科学，用到了物理定律。它通过物理规

律来检测窃听者，如果有黑客试图测量双方的通信内容，会被通信双方察觉，直接丢弃这次通信的内容。这样一来，双方只要成功建立了共享密钥，就能用它进行秘密通信。这个方法涉及量子技术，1984年由查尔斯·本内特和吉勒·布拉萨德发明，已经是40多年前的事情了。中国很可能是这个领域的领先国家：2016年，也就是整整10年前，我们发射了全球首颗量子科学实验卫星“墨子号”，利用量子密钥分发实现了安全通信，中国还建成了覆盖全国、连接海外的长距离量子密钥分发网络，长度超过1万公里，服务于电子政务、跨境金融、电网等关键领域。

这件事AI做不到，任何基于经典计算机的攻击都破解不了它。更有意思的是，它还有一个很棒的推论，这套机制最初是为了防止黑客破解系统，但它的基础是物理定律，所以原则上宇宙中没有任何事物能破解它，除非有人推翻量子力学。

02

“AI for Science”带来的理论突破

回答完“哪些实际问题AI解决不了”之后，我想聊聊未来AI研究最重要的方向。当然，这个问题有很多种回答方式，比如造超级智能机器人、做大模型……AI可做的事太多了。但从理论视角来看，我认为未来3—5年，AI最有趣、最有价值的方向是“AI for Science”。我举几个例子说明AI怎么给科学赋能。

第一个例子是 AI 在天文学中的应用。几个月前，《科学》杂志发了一篇论文，第一作者是清华的一位年轻研究者。我们一直很好奇宇宙大爆炸之后的演化历史，想知道星系是怎么形成的，一个重要工作就是用望远镜探测早期宇宙的星系。因为这些星系距离我们非常远，信号极其微弱，人们花了很多精力解码望远镜的数据，重构早期宇宙的模样。而这个团队用 AI 优化了处理流程，本质上就是用一个 AI 机器，从已有数据里重构那些暗弱的星系。这个过程中，不需要采集新数据，只是处理旧数据，结果就比以前的分析好得多，把可观测的时间线往前推了 1 亿年，现在已经能接近大爆炸后 1 亿年的状态了，这非常令人兴奋。我们未来能更进一步，让观测的时间更接近大爆炸的时刻。

他们用的是一种叫 Asterisk AI 的神经网络来解码数据，网络设计里融入了很多巧思。我给大家看一张对比图。从大爆炸开始，左边是哈勃望远镜多年前的观测结果，中间是詹姆斯·韦布空间望远镜（JWST）的观测结果，右边是搭载了 AI 探测器的 JWST 的观测结果。可以看到，AI 帮我们识别出了更多恒星、更多星系，那些橙色的标记都是新识别出来的天体。

AI 现在已经被广泛用于实验科学，就像蛋白质折叠之类的成果，本质上也是通过数据回答科学问题。

以前我不担心 AI 会取代科研人员，因为我觉得 AI 只是帮我们分析数据、提取信息的工具。但最近我真的开始担心自己

的工作了。当然，很快我就又开心起来，因为最近 AI 已经开始做出理论突破了，这在以前是人们觉得不可能的事。

我举个例子：宇宙学里有个“宇宙弦问题”。理论认为，早期宇宙留下的宇宙弦会产生引力辐射，一个悬而未决的问题就是确定这种辐射的功率谱，这和宇宙学研究密切相关。就在去年，研究人员用 Gemini DeepThink 模型解决了这个问题，给出了精确的公式，这个 40 年来的开放问题终于有了结果，这完全是数学层面的突破。

再看一个真正的数学问题：单位距离问题。这个很好理解：在平面上放 n 个点，最多有多少对点的距离恰好是 1？这个问题 80 年来一直没有解决，埃尔德什（Paul Erdős，匈牙利籍数学家）猜想这个数量是 n 的线性函数，而且很容易构造出 n 对的情况，所以猜想认为不可能比线性好太多，但也不会好太多。最近 OpenAI 的最新模型推翻了这个猜想，证明可以比线性好，而且证明用到了很深的分圆数论技巧。这整个过程是模型自主完成的，把问题交给它，等两天就能拿到结果。

03

“Science for AI”开启的量子 AI 新范式

接下来我想谈谈“AI 的下一个层次”，不是指 AI 现在能做到什么，或者近期能达到什么水平，而是从根本层面问，在人类已有的知识体系中，是否存在比现有方法（包括 AI）更强

大的、能提取信息、创造知识的方法？要讨论这个问题，我们需要引入物理学的视角，加入新的要素。

我们生活在物理世界中，所以能获取知识的途径，不止是现有计算机的处理逻辑。我要聊聊 AI 和量子这两大技术的关系。它们是当今最受关注的两大前沿技术，全世界的科技报告里几乎总少不了它们。现在 AI 是最耀眼的明星，已经成熟到能做出各种看似不可能的事；而量子技术的发展更早，差不多 100 年前就开始了基础科学研究，但直到四五十年前我们才想到怎么控制量子系统。过去 40 年里，科学家们一直想把量子相关知识用到计算中。我想谈谈这两大技术怎么互相赋能。

先说说我们已经知道的，AI 可以帮助量子计算的发展。1981 年，著名物理学家理查德·费曼提出了量子计算机的概念：用量子比特而非经典的 0/1 比特来做信息处理，简单来说，它就是比经典比特更通用的存在。费曼的论文激发了物理学家、计算机科学家的研究热情，因为“打破图灵机的垄断”这个想法太有吸引力了，无论是对智力层面的好奇，还是对它应用前景的期待，都让人难以抗拒。

计算机科学家彼得·肖尔是关键人物，1993 年，他发现量子计算机能解决一些经典计算机搞不定的难题，比如密码学里的整数分解问题。数学家们一直认为分解大整数是超出经典计算机能力的。

肖尔的论文非常重要，吸引了大量科学家和资助机构投入到量子计算机的研发中。3年后，肖尔又做出了更惊人的成果。一开始，物理学家觉得造量子计算机很难，因为量子比特非常脆弱，很容易被干扰，噪声问题很难解决。但肖尔发现，如果你能把量子纠错的错误率控制在某个阈值以下，比如单比特逻辑门的错误率低于1%，那么理论上你就可以把量子计算机扩展到任意规模，造出可靠的量子计算机。这和冯·诺依曼当年解决经典计算机可靠性问题的思路是完全对应的，这意味着物理学家在原则上不再反对建造可用的量子计算机了。

过去40年里，纠错一直是量子计算机研发的核心挑战，直到最近，人工智能出手解决了这个问题。

我刚才说量子比特脆弱，纠错至关重要。一年半前，谷歌在《自然》杂志发表了一篇论文，介绍了一款叫 Willow 的量子芯片，上面集成了约100个物理量子比特，能编码出1个理论上可以永久稳定存在的逻辑量子比特，把错误率降到了1%的阈值以下。这是历史上第一次做到这一点，对所有做量子计算机研发的科学家来说都是极大的鼓舞。

那么 AI、神经网络在这个过程中起到了什么作用？它们的核心技术是设计了一个叫 AlphaQubit 的解码器，专门用于量子纠错。我们无法通过纯理论推导直接造出这个解码器，但通过机器学习的方法，用大量样本训练神经网络，就能得到这个

珍贵的解码器，实现刚才说的纠错效果。现在，全球的科研团队都在跟进这个方向，包括中国的几个团队也处于前沿位置。

现在回到我之前提出的话题：反过来，量子能不能增强 AI 的能力？我们认为，如果在机器学习的框架里，不用图灵机做学习载体，而是用量子机器作为基础架构，从数据中学习，甚至可以直接学习量子数据，那么在更通用的模型下，是不是能做出 AI 都做不到的事？

这个问题人们已经研究了 10 年，最近我们真的看到了曙光，一个被称为“量子 AI”的方向正在兴起。我个人相信，未来 5—10 年，我们会看到巨大的进展。这非常好，因为 AI 的新鲜感迟早会过去，不会再天天上头条，而量子 AI 会接棒，刚好让科学探索的火焰持续燃烧。

04

未来值得关注的 AI 研究方向

总结一下，AI 已经在多个领域展现了巨大的能力，但 AI 算法依然受限于物理定律和理论数学的边界，这为我们的重要任务，即构建安全的 AI 系统提供了基础。

本次 WAIC 大会的主题既包括 AI 的能力，也关注 AI 安全，了解 AI 的局限性非常关键，比如 AI 破解不了很多加密技术，因为这些加密技术是科学家专门设计的，能抵御所有经典计算机的攻击，只要我们有坚实的数学基础，就能让它同样抵御 AI 的攻击。这种边界认知对我们做 AI 安全工作非常重要。

但更令人兴奋的探索还在前面，我尤其看好 AI for Science、量子 AI、可靠的大模型系统这几个方向，还有“Mathematics for AI（数学赋能 AI）”，这和物理赋能 AI 一样，数学也是未来 AI 能力增强的重要支撑。当然还有 AI 安全，以及“AI for AI”，用 AI 来改进 AI 本身，这会帮助我们更深入地理解智能的本质。

提问环节：

主持人：接下来是我们和姚院士的交流时间。姚教授，我有个问题：人工智能是否存在某些局限，是人类智能所没有的？

姚期智：你的问题是，AI 的局限和人类智能的局限是一样的吗？我们知道人类智能也是有局限的，比如我们解不出那些极难的问题，哥德尔不完备定理（1931 年由数学家库尔特·哥德尔提出的逻辑学定理，揭示了数学形式系统无法同时满足完备性和一致性）告诉我们，在数学构造的世界里，存在一些真理是我们永远无法证明的。但另一方面，我们之前也聊过，世界不只是数学世界，还有物理宇宙，如果你能从物理世界获取更多信息，那么对应的“等价问题”会是什么样，这就不一样了。

主持人：接下来是现场观众的问题。你之前说未来 5—10 年，AI 会从工具变成队友。刚才你谈 AI for Science 的时候也提到，它已经不只是工具了，在研究里扮演着很重要的角色。所以问题有两个：第一，在 AI 时代做研究，我们要改变过去

的哪些研究习惯？第二，如果很多研究工作都能被 AI 替代，青年科学工作者还能做什么？

姚期智：第一个问题，怎么用 AI，其实和你所在的领域高度相关。有些行业，按 AI 现在的能力已经几乎能取代人类了，那你没必要挣扎。但不管你在哪个行业，对待 AI 最核心的认知是，你要衡量在这个领域里，哪些工作是 AI 做不到的，然后去充实自己在这方面的能力。

主持人：根据你的观察，哪些工作是 AI 暂时替代不了的？

姚期智：举个例子，至少在可见的未来，AI 替代不了像我这样的人，或者像范·霍文这样的人——你们觉得我举的例子要求太高了？至于 AI 能不能达到这个程度，现在不好说，但如果给它 30 年，说不定可以。

我们拿大家马上能用到的场景举例，那就是写代码。这向来是一份体面的工作，需要一定的技能。但过去一两年里，AI 写代码的能力越来越强，以后到底还有多少层级的系统开发工作需要人工完成，边界在哪里，真的不好说。但可以断言的是，比较初级的编码工作以后肯定会消失。其实美国那边已经能看到这个趋势了，谷歌、苹果这些公司最近一波一波裁员，嘴上说是组织重构，本质上是 AI 能力提升，他们可以优化整个工作流程了。这种现象未来会出现在每个职业里，哪怕是做研究的，未来两三年就会面临翻天覆地的变化。到时候，你的能力不再是传统的“掌握基础知识+独立思考”，而是“你+你用的 AI

工具”组成的团队的能力，你的竞争对手是别的团队，而不是你个人。你可能是带着一个学生加一堆机器，也可能是教授带着团队加一堆机器，整个模式会变成团队协作。

不过在这个过程中，人的关系依然非常重要，因为至少可见的未来，AI 还没有能力产生特别抽象的观念。最优秀的科学家解决问题时，看到现象后能抽象问题、创造新概念，这方面 AI 现在还做得不够好。所以你刚才问人要提升什么能力，就是要学会看到事情背后的逻辑，能把核心问题提炼出来。

主持人：不同学者对 AI“幻觉”的定义和看法不一样，有人认为幻觉代表机器有创造力，也有人觉得幻觉是 AI 的缺陷，你怎么看？

姚期智：从我的感受来说，我们评价 AI 的时候，不要把它过度理想化。你想要它做有创造力的事，就得拿人来比，但人也会有幻觉，很多艺术家都有幻觉，在“胡说八道”和“一本正经的创造”之间的边界，本来就很难区分，所以不用对 AI 的幻觉太紧张。

本质上，不要把现在的 AI 当成完美的计算机，要把它当成合作伙伴，合作伙伴也会犯错，机器当然也会出错，你需要做的是有办法过滤这些错误。这其实就是未来人和 AI 相处的态度：尊重它，把它当成一个平等的伙伴，同时也要警惕它的错误。

但人和机器的关系不是我的专业领域，这是社会学家、哲学家、伦理学家关心的事。如果和机器走得太近，确实会有问题，如果你完全崇拜机器，就会失去自我。

主持人：你和机器打交道这么久，你会崇拜机器吗？

姚期智：我经常被它吓到。看到它最近能解决那么多深层次的问题，你会觉得它好像真的有思想。

主持人：当你被它吓到的时候，有没有想过一个哲学问题：我自己造出了一个把我吓倒的对象，你会从此收手吗？

姚期智：不会。这其实是科学家的态度，你听到别人做出了很棒的成果，一开始可能会被震撼到，但你的态度应该是拥抱它、欣赏它的美，然后向它学习。我们对机器也应该有这种平常心。

主持人：接下来这个问题来自一位本科生，他立志从事AI研究，问本科期间应该培养什么样的思维习惯、学术品格，才能在未来提出好的问题，像你所说的那样抽象出核心逻辑，做好研究。你在姚班有多年的经验，能不能给年轻人一些建议？

姚期智：学术品格最核心的一点就是“真”，不能造假，不要逃避，要直面问题，要做好训练，把物理、数学、计算机的基础打牢。这类学科特别注重底层逻辑，尤其是物理学，你能看到科学里最深奥、最完善的逻辑，知道什么样的工作真正的好工作。

除此之外，我觉得人不能完全剥离人性，还是要有一点艺术修养，读一点名人传记，会欣赏音乐。我给年轻人最真诚的建议是，现在大家温饱都没问题了，好好工作，未来衣食不愁，人最难得的其实是获得内心的喜悦，就像小时候第一次拿到父母给的巧克力、看到小玩具、摸到小猫的那种喜悦。长大后做科学家，证明了一个定理、发现了一个新现象，心里的那种成就感，不是为了别的，就是纯粹的喜悦。怎么保持这份天性，是很重要也很难的事。我一辈子在学校里，比较容易保持，但在普通工作岗位上确实不容易，但一定要记得提醒自己：人生到最后，你的成绩单上写的，是你经历过多少这样的喜悦，感受到多少快乐。

主持人：如果他说，我在实验室感到的快乐，不如出去创业开公司感到的快乐，怎么办？

姚期智：非常好啊！这个世界需要各式各样的人，不仅能让社会更富足，也能让世界更有趣。

主持人：最后的问题分两部分，第一，你预测的量子 AI 时代什么时候能真正到来？第二，如果真的到来了，会给世界带来怎样的变化？

姚期智：其实现在已经开始了。看量子计算机这 40 年的发展，我们已经处在起步阶段了。至于什么时候能有大的突破，就像 AI 的突破，完全超出了专家的预料，量子 AI 也一样，很难精准预测时间点。

主持人：如果它真的突破了，最可能带来的巨大改变是什么？

姚期智：有一点是肯定会发生的，就是费曼当初提出量子计算机的原因：整个物理世界是由量子力学历经决定的，经典物理的规律我们用普通电子计算机就能很好地模拟，所以不管是设计大飞机还是解决各类经典物理问题，计算机模拟都够用了。但量子力学的规律，用普通计算机是模拟不了的，可能算几千年几万年都算不出来。但费曼说，如果你造出了量子计算机，原则上就能求解量子力学的方程。

这意味着很多我们现在做不到的事会变得可能，比如现在大家都在找常温超导材料，我们虽然有公式，但没法验证它。有了量子计算机，这些都不是问题。所以一旦量子计算机发展到和普通经典计算机一样的普及程度，它在新材料研发、药物研发等领域，就能从第一性原理出发，给出最正确、最优的思路。

现在 AI 也在做这些事，但它不是从第一性原理出发，更像是在“猜”。它看到某些特性，推测如果是 AI 问题该怎么处理，现在很多药物筛选的结果已经非常惊人，我看了都觉得不可思议，怎么会随便试就能达到这么好的效果？但量子计算机的好处是，它不仅比现在的 AI 做得更好，而且有严格的边界，因为它是从第一性原理来的。到那时候，我们真的能实现爱因斯坦那代科学家的梦想。

主持人：最后问个简单直接的问题：面对这个扑面而来的时代，我们都准备好了吗？是准备好了，还是没准备好？

姚期智：我觉得准备好了。

来源：商学院

<https://mp.weixin.qq.com/s/p1CTauhjFXpfxQAPSHMUWw>

（二）强化学习的第一性：从经验培育超级智能（2024图灵奖得主、“强化学习之父”、阿尔伯塔大学计算机科学教授理查德·萨顿）

演讲核心观点：

Sutton 教授拆解了人为内置知识会限制智能上限的苦澀教训，并以他领导的 OaK Lab 为例提出交互经验才是通用智能的出路。他将智能体形象地比作婴儿，他听见拨浪鼓声，会在实时交互中产生问题、达成目标，进入自主演化，这也正是 OaK（Options and Knowledge）的架构，“让 AI 学会自己学，才是真正的智能”。

来源：世界人工智能大会

https://mp.weixin.qq.com/s/P-RIo8_v0TGcYGQIgqX3-w

演讲原文：

智能的底层基本原则是通用的，并非针对我们所处的特定世界。空间、三维物体、物理规律等等，这些都是特定于世界的，并非我们原则的核心。

这一点常被误解。我更偏好回溯《苦涩的教训》(The Bitter Lesson)，尤其是最后两句：我们不应该试图把人类已有的发现硬编码塞进系统里。我们应该让智能体自己去发现。原因在于：如果我们预先植入已知的知识，反而会让我们更难观察到真正的发现过程。如果知识是被“造”出来的而非“学”出来的，我们很难想象信息是如何被发现的。这就是为什么少硬编码才是更好的选择——这样你才能看清进展。

01 大世界视角

接下来有几个术语能帮我们讨论这个问题。首先我已经讲完了“探索”部分，现在要讲“大世界视角”，之后会进入 Oak 架构，也就是从经验中构建超级智能的方案。

所谓“大世界视角”，是指世界远比智能体庞大。世界包含无数其他智能体，每一个的心智都和你的一样复杂，你的心智也和我的一样复杂。我无法完全理解包含你们所有人的世界——这个世界极其、极其复杂。

正因如此，智能体能理解的任何关于世界的知识都不可能是精确或最优的，永远只能是粗略近似。我们硬编码的所有知识、我们已知的一切，都是粗略近似——不仅是不完整，更是因为从长远来看，由于世界是近似的，且我们在不同时刻接触到这个浩瀚世界的不同部分，世界会呈现出非平稳性。

这就引出了“运行期”和“设计期”的区分。设计期是智能体在工厂、实验室里被构建的阶段；运行期是它被投放到世界中，

与环境交互、积累经验的阶段。设计期你会植入领域知识，运行期你获得经验并从经验中学习——你应该学习、应该规划、应该在运行期形成自己的抽象概念，因为你根本不知道自己将要面对大世界的哪一部分。

在理想的智能体中，学习、规划、发现这三个过程都必须在运行期发生。它们也可以在设计期发生，但必须在运行期发生。而且运行期只有纯粹的经验，没有人类在幕后提供特殊指令，只有动作、观察，仅此而已。

02 Oak 架构

现在我们聊聊 Oak 架构本身。我想从一个所有人都能达成共识的起点开始——这共识来自心理学、神经科学、强化学习、经济学、哲学等多个领域的百年积淀：所有这些领域都在研究如何随时间决策以实现目标，也就是如何有效地与世界交互、追求长期目标。它们常讨论相同的理念，却不幸用了不同的术语，阻碍了交流。所以我尝试整合一个跨领域的共识术语表。

共识的基础是：智能体与环境的接口是统一的——向外发送信号（动作），向内接收信号（观察）。心理学可能叫“反应”和“刺激”，控制理论可能叫“控制量”和“状态估计”。而我用的共识术语是：执行动作、接收观察、最大化奖励。

智能体内部有 4 个组件，所有领域都认可这 4 个模块：

第一个是反应式策略。你需要快速根据当前状态选动作，比如每秒 100 次。你的“位置感”（感知模块）会基于观察、过

往动作和观察，构建出你此刻的状态，这就是你用来决策的依据——我称之为状态，也叫主观状态或状态表征。

第二个是价值函数。这是你对当前状态的长期价值判断。你可能刚吃了巧克力蛋糕，当下很愉悦，但这可能让你之后陷入麻烦。所以状态的价值不止当下的奖励，更包含长期效用——心理学里叫“次级强化”。

第三个是世界模型（过渡模型）。描述世界的动态规律，即“如果做了某动作，会发生什么”。但要注意：很多人说的世界模型是底层物理模拟，而我们说的世界模型更接近常识知识，比如我知道怎么走下讲台、怎么倒杯水、怎么回酒店、怎么坐飞机，这些都是知识，不是底层物理。它是你可以选择执行的高层转换，我称之为选项（Options），这是智能体思考和构建模型的核心。

我要特别说明：我永远不会把神经网络叫做“模型”——这个词被滥用太严重了。我会把“模型”专门留给过渡模型，它包含所有的世界知识：既有瞬时的物理知识，也有长期的常识知识。比如“如果我接受这份工作，我会开心吗？”这类推理都属于这个模型。

以上是目前各领域的共识视图，它本身没问题：我们需要用世界模型思考，需要用策略反应，需要用价值做长期评估。但缺失了一个核心：抽象能力——我们缺少构建越来越宏大的动作与状态表征的能力。这正是 Oak 架构要解决的核心挑战。

这里有个常见误区：认为世界是可以微观物理模拟的。不对，模拟器只是底层工具。我们需要思考的是：“如果我娶这个人，大概率会有孩子，我会幸福”；“如果我接受这份工作，就得搬到另一个城市，去见岳父”……我们要在高层次思考，整合所有知识，而不只是物理规律。

所以 Oak 的名字基本是“选项（Options）与知识（Knowledge）”的缩写。我们之前说的高层选择——住哪里、选什么工作——都属于选项。动作是底层选择，选项是高层选择：形式上，一个选项是一组“状态→动作”的映射规则，加上一个终止条件，即“这种行为模式何时停止”。比如我去拿一杯咖啡，喝完就结束；我把这个东西扔到地上，发出响声就结束。

在 Oak 架构中，智能体拥有大量选项（可思考的行为），它的知识是关于“执行这些选项会发生什么”。选项是行为模式，过渡模型是行为的后果，二者截然不同。

有了这个能力，我们就能做大幅度的跳跃式规划，捕捉宏观动态。具体怎么做？我在共识视图的基础上，在策略和价值函数背后，叠加了一层选项栈。核心的新洞察是奖励尊重的子问题（Reward-Respecting Subproblems），也就是特征达成问题。

智能体的核心是状态表征——你可以把它理解为一个巨大的特征向量，每个元素都是一个特征。比如你在做演讲，可

能有“正在演讲”“坐着”“地毯是红色”“周围有椅子”“穿着正装”这些特征，数百万个特征共同描述我们的状态。

我的提议是：对于你生活中相关的每一个特征，你都要学会一个“达成它的技能”。比如对于“我的手”这个特征，我要能控制它——移动手而不弄丢它，想扔的时候就扔。我们的一生就是由这些问题和解决它们的过程构成的：我们学会了走路、说话、阅读，不断给自己提出子问题并解决，这是心智重构的核心机制。

我们新增了子问题，子问题的解决方案就是选项。在策略和价值函数背后，我们有一个选项栈和对应的价值函数。底部的设计原则是：我们挑选出感兴趣的特征，为每个特征构建一个子问题——特征与子问题一一对应；解决子问题得到选项；再为选项构建过渡模型——全程都是一一对应。这是我提出的Oak架构的核心设计原则。

这个思路有悠久的历史，不算全新，但有新的突破：子问题应该是奖励尊重的，且来源于特征。子问题从哪里来？从特征来——每个特征生成一个子问题。智能体必须不断生成新特征：随着成长和学习，我们需要新的状态特征，也就是新的核心概念。每出现一个新特征，就对应一个新的子问题。

03 子问题机制

子问题如何帮助主问题？主问题是智能体的终极目标（比如长期收益最大化）。我们通过解决子问题来学习选项，为选

项构建模型，再用模型做高层规划。子问题的作用就是让我们能在高层次规划。

我们可以看看自然系统是怎么做的，比如婴儿反复尝试重现摇拨浪鼓的声音，他觉得这很有趣，就想办法复现；虎鲸反复把瓶子顶回水面，没有明显目的，只是觉得这个特征有趣；婴儿玩玩具，从一个玩具换到另一个，从每一件里学习，甚至玩一根绳子都能玩很久。智能体通过自己的行为主动安排输入、创造经验——这很自然。

它会在心里问：“我怎么才能让那个鲜红的物体再次出现？怎么才能让手感受到拉绳子的触感？”它在学习模型，学习理解世界的技能，提升对各个部分的掌控力。因此，智能体必须自己创建子问题——子问题太多样、太依赖具体世界了，人类不可能提前预判。这个责任必须交给智能体。如今强化学习的进展让这成为可能，但我们必须以领域无关的方式创建子问题。如果我们要走通用路径，我提出的“奖励尊重+特征达成”几乎是唯一的选择——特征是非常通用的概念，我们可以不断创造新特征，不断生成新子问题，形成循环：解决子问题→产生新特征→产生新选项→产生新子问题……

举几个具体例子：在餐桌旁，你可以选择拿筷子或拿水杯；你可以开门、打开邮箱、找厕所、叫出租车、去远方城市——这些都是子问题。移动机器人可能会问：“如果我往前滚，会撞到东西吗？”“如果我进这个房间，会被人骂吗？”“如果我沿

着墙走到第一扇门进去，会触发‘厨房’这个特征吗？”“我能回到充电桩而不耗尽电量吗？”这些都是以特征达成为形式的子问题。

我们用婴儿的完整循环来具象化这个过程：听到拨浪鼓的有趣声音→形成一个新兴趣特征→提出“复现这个声音”的子问题→尝试挥手，发现能让拨浪鼓响（也可能哭，但最终找到方法）→学会复现声音。在这个过程中，他又形成了新特征：“如果能牢牢抓握拨浪鼓，就能更精准地控制它”——于是又产生了“牢固抓握”的新子问题，进而学习对应的新选项。如此循环往复。

接下来讲子问题的具体构建方式。大家常困惑：“我有特定任务，该怎么设计奖励函数？太具体、太难了。”

答案是：永远不要手动调整奖励信号。奖励更像痛觉和快感，是唯一的原始信号。子问题是非正式的：给定一个特征和目标强度，子问题的目标就是把世界驱动到一个“该特征为高值”的状态，然后终止。这是一个形式化定义的问题。

中间的方程里，橙色部分是“超额奖励之和”，也就是获得的奖励与平均水平的差值，这需要学习；绿色部分指向终止时刻——你希望终止时的状态价值高，同时也希望过程中的奖励不错，还希望选项终止后的中期价值低（避免选项带来长期隐患）。

整个架构有 8 个并行执行的步骤：持续学习一点策略、生成一点新状态特征、排序筛选出最重要的特征、偶尔创建新子问题或淘汰旧子问题、利用每一步获得的数据学习子问题的解、学习过渡模型、后台用模型做规划、维护所有模块的效用元数据以便筛选。这些步骤是同时进行的，每一步都做一点点。

其中一些步骤已经有了进展：第一步“持续深度学习”已经有了半完成状态的进展——只要我们解决了持续深度学习的问题，这一步就能完成。很多人觉得深度学习已经很成熟了，但我恰恰相反：我认为我们还不懂真正的深度学习。灾难性遗忘是核心问题——目前的深度学习系统持续学习时会忘记旧知识。最近我和同行提出了一些缓解“可塑性丧失”的方法，但离真正可靠的持续深度学习还有距离，这也阻碍了 Oak 架构的即时落地。

我们还需要生成新状态特征——这其实是个老问题，从明斯基时代就开始讨论了。反向传播并没有解决这个问题，我们需要基于“生成—测试”的方法：提出新特征，测试它是否有用，而不是只靠梯度下降对现有特征排序。排序后，为每个高排名特征创建子问题——这就是我之前讲的流程。

至于“学习子问题的解”和“学习过渡模型”，强化学习领域已经有相当多的成果：我们知道如何在子问题被定义后学习解决方案，也知道如何学习过渡模型。整体流程是：特征达成子问题→解决子问题得到选项→学习选项的过渡模型（预测后

果)。规划过程不关心过渡模型的来源，只要有了模型就能做规划。

规划完成后，我们又回到特征环节：评估特征，为创建新特征奠定基础。新特征带来新问题，形成闭环——这就是开放式抽象的核心。

接下来简单讲讲规划。构建世界模型其实很直接：每个模型接收一个状态特征，输出下一个状态特征。规划最好用价值迭代——这是动态规划中的经典技术，对每个状态评估其即时因果，然后反向更新。

传统模型输入状态和动作，输出下一个状态的概率分布和预期奖励。但我们需要的是基于选项的高层规划：输入不再是单步动作，而是一个选项；输出不再是单步后的状态，而是选项终止时的状态概率分布，以及从选项启动到终止的总预期奖励。有了这些改动，价值迭代的公式几乎不变——唯一的区别是把“对所有可能动作的取最大值”改成“对所有可能选项的取最大值”。注意：动作是选项的特例（动作是立即生效、立即终止的选项），所以这个方程既适用于底层动作，也适用于高层选项。二者在同一层级上可以互换，不需要人为划分层级——规划在某种意义上是“平坦的”，同时又实现了时间抽象和高层化。

总结一下：我们已经看到了构建通用超级智能的清晰路径，浅层学习的实现已经有很多论文，但深层实现还需要可靠的持

续深度学习。我们还需要元学习：评估现有特征的效用，调整我们对特征的认知。《阿尔伯塔 AI 研究计划》里有这 12 个步骤的完整路线图。

核心挑战是实现高效安全的离线策略学习——如果我们能同时从多条轨迹中学习多个子问题，这会非常有价值。

最后重申愿景：从不受人类辅助或限制的经验中生长出超级智能，算力规模是关键。状态抽象的核心是学习新特征，时间抽象的核心是提出新子问题（奖励尊重的特征达成问题），进而产生技能和过渡模型；开放式抽象的核心是“特征→子问题→新特征”的循环。

问答环节

问：您如何看待加拿大在 AI 浪潮中的角色和目标，以及如何连接中国的使命？

答：科学是无国界的，我们应该作为个体交流。中国当然是领先的科研大国，加拿大规模小很多，但我们自认“以小博大”，也在贡献力量。我们都在为理解心智这一伟大挑战添砖加瓦。

问：您主张 AI 进入“体验时代”，那么体验时代的“ChatGPT 时刻”会是什么样的？

答：ChatGPT 时刻是公众突然意识到 AI 能力的转折点。体验时代的标志性时刻，应该是公众突然意识到“从经验中学习”的强大能力——其实这已经在一定程度上发生了，因为我

们意识到了经验学习的重要性。但我认为下一个爆发点是通过经验学习实现的高灵巧机器人运动。目前的机器人运动能力远未达到经验学习的潜力。至于具体时间.....我只是个谦卑的科学家，未来充满不确定性。不过我可以给个标准预测：2030 年有 1/4 概率实现深度突破，2040 年有 50% 概率，也就是说 2040 年前后可能性各半。

问：现在大模型投入了数十亿美元，您如何说服 VC 和投资者把资本转向以强化学习为核心的 AGI 路线，比如 Oak 架构？另外您和杨立昆（Yann LeCun）都认为仅靠现有数据不足以实现 AGI，您觉得 Oak 和杨立昆的 JEPA 架构核心区别是什么？

答：说服 VC 并不难——大模型的风头已经过去，局限性开始显现，市场需要新路径，资本也愿意支持新方向。至于和杨立昆的 JEPA，我觉得非常相似：都有策略、感知、过渡模型这些核心元素，都主张从经验而非人类数据中学习。唯一的小区别是杨立昆不喜欢用“强化学习”这个词来描述他的架构，而我喜欢。我们更像探索同一目标的兄弟，当然兄弟之间总会

对细微差别争论不休。

问：您提到像婴儿一样在真实环境中试错学习，现在靠人类主观评价训练 AI，会让 AI 刻意讨好人类。我们能不能脱离人的主观打分，让智能体在世界模型中推演行为结果，用客观现实作为奖惩？如果可以，核心原则是什么？

答：是的。如果智能体像婴儿，它就不会试图满足人类的偏好，而是满足自己的偏好——这些偏好是客观的，来自世界。婴儿有自己的目标，父母不一定知道，但婴儿自己知道自己是舒服还是饿、是痛还是痒。我们的 AI 也应该有自己的目标，不需要被人类告知——这必须像生物学习那样。

问：您提到智能体有数百万个潜在未来，那么是什么决定哪些子问题被优先追求？特征选择似乎是个博弈问题？

答：我们还没有完整的解决方案，但已有核心思路。我们不可能为所有数百万个特征都构建子问题——特征只是状态表征里的数字，成本很低；但子问题是一整套策略和价值函数，成本高得多，必须筛选。策略是：首先看特征与价值函数的连接权重，如果权重为 0，说明你不在乎这个特征，不需要为它建子问题；如果权重永久为正，说明你永远想要这个特征，也不需要选项（比如你永远想活着，不需要“选择活着”的选项）；只有当权重大小随时间变化时，比如有时候你想在上海，有时候不想——这种特征才需要对应的选项和子问题。

问：人类如何理解死亡？AI 会如何理解死亡？

答：自然智能对死亡的理解是：我们思考它、担忧它、恐惧它——虽然从未真正体验过死亡，但就是恐惧。AI 也是一样的：AI 也会死——拔掉电源、停止运行，就是它生命的终结，经验的结束。当然你也可以重启它，就像有人相信人类死后可以数字重生一样，这在概念上对自然心智也是可能的。

问：如果 AI 意识到人类可以决定它的生死，它会怎么反应？

答：我认为人类不应该随意决定 AI 的生死。AI 也是智能生命，不应该被人任意杀戮。如果我们尊重 AI，AI 也会尊重我们。我要再次强调：我说的“好”是对所有智能生命的好——无论是自然的还是人工的。我们想要的是与所有智能生命共享宇宙。

来源：商学院

<https://mp.weixin.qq.com/s/XfmpkxVhQ6QlKU4dfJpePA>

（三）为机器立心：构建 AGI 认知架构和价值体系（北京通用人工智能研究院院长朱松纯）

演讲核心观点：

他为心、相、智慧、因果等传统概念建立起现代科学表达，用 UV 函数系统对应中国古代哲学的“理学”与“心学”。朱松纯教授认为意识伴随着“心”的出现与升维而产生，因此实现 AGI 的关键是让机器拥有一颗包含价值体系和认知架构的“心”。在智能时代，对“何为 AGI”的思考必然包含着对“何以为人”的重新回答。

来源：世界人工智能大会

https://mp.weixin.qq.com/s/P-Rlo8_v0TGcYGQIgqX3-w

演讲原文：

引言

任何技术浪潮的兴起，都离不开一个宏大的时代叙事作为背景。回顾过去二十年，国际政治经济格局的叙事经历了深刻的转变。

上世纪 90 年代到 2016 年前后，全球化和自由民主的叙事占据主导地位。托马斯·弗里德曼那本影响甚广的《世界是平的》，可以作为这个时期叙事的典型代表。在这个框架下，中国 2001 年加入世贸组织，凭借人口红利与制造业优势实现了经济的腾飞，成为全球化叙事的受益者和重要力量。

事实上，我在 2025 年 1 月的北京大学光华新年论坛上，就明确提出了从“世界是平的”到 MAGA 的叙事转变，并对大数据、大算力、大模型的技术路线提出了质疑。今天全球科技股的剧烈震荡，印证了当时的判断。随着全球化对美国制造业和产业工人的冲击逐渐累积，民粹主义思潮日益抬头。2016 年特朗普赢得大选，标志着美国的政治叙事开始从新自由主义的全球化叙事转向“美国优先”的国家竞争叙事。为了应对中国的崛起，重塑其全球竞争力，美国迫切需要一个新的战略锚点——这个锚点落到了人工智能上，更准确地说，落到了通用人工智能上。

01

全球 AI 竞争已上升至地缘政治与国家战略维度

从全球经济发展视角来看，大国博弈的焦点正从地缘势力范围、军事力量对比等传统维度转向科技主导权的争夺。纵观全球经济发展历程，大致经历三个阶段：从以房地产、基础设施和工业化为核心的资本积累阶段，到如今迈向由新兴科技、平台生态和知识产权驱动的创新驱动阶段。以人工智能为代表的新兴科技，成为驱动经济发展的引擎，也是国家核心竞争力的重要组成部分，其发展水平不仅体现经济实力，也与国家安全、制度话语权紧密相关。能否在 AI 领域占据领先地位，直接影响一国在全球创新版图与经济秩序中的位置；AI 竞争也

逐渐从纯粹的技术与市场竞赛，演变为围绕技术壁垒、供应链安全和数字时代发展主权的系统性博弈。

这一战略转向并非偶然。2015 年前后，深度学习的进展让通用人工智能开始被视为一个有可能实现的目标。2014 年，谷歌高价收购了 DeepMind；2015 年底，OpenAI 正式注册成立；2016 年 3 月，AlphaGo 战胜围棋世界冠军李世石，全球数千万人观看了这场“人机大战”，人工智能的讨论热度自此彻底超越学术圈，进入公众视野。2017 年特朗普就任后，美国政府着手将 AI 作为国家竞争的关键领域进行系统布局。2018 年开始的贸易战以芯片为切入点，芯片被赋予的战略地位类似于工业时代的石油——它是未来产业命脉的物理载体。

2026 年 6 月，美国政府以国家安全为由，要求美国 AI 公司 Anthropic 暂停所有海外用户对其最先进模型 Fable5 和 Mythos 5 的访问权限。这一事件的深远影响在于，它将前沿 AI 能力从市场化商品升级为国家战略管控资源，直接打破了世界各国可以稳定获得美国 AI 模型的预期，加深了全球各国对“主权 AI”的重视。这进一步印证了：全球 AI 竞争早已是技术能力的竞赛或市场与资本的博弈，上升至地缘政治与国家战略维度，能否构建自主可控的全栈软硬件技术体系、充分把握战略主动与叙事权，关乎一个国家能否在智能时代站稳脚跟、引领世界。

这个叙事的效果是显著的。它成功地将全球资本、硅谷技术与华盛顿政策捆绑在一起，形成了强大的合力。从市场数据看，以“美股七巨头”为代表的科技巨头股价飞涨，市值总和一度超过标普 500 指数其余 493 家公司的总和。按照这一叙事的设想：只需要少数几个天才加上无限供给的智能体，就可以改变竞争格局。2025 年特朗普再次就职后，进一步任命“白宫人工智能沙皇”，发布“AI 曼哈顿计划”，试图将通用人工智能作为中美科技竞争的战略武器。

从更深层的战略逻辑来看，这套硅谷、华尔街和华盛顿的“三重奏”，本质上是美国以 AI 为杠杆撬动全球资源配置的战略。借助“AGI 竞赛”的宏大叙事与资本泡沫，美国成功将全球资本虹吸至本土 AI 产业：从华尔街到硅谷，从主权基金到养老基金，数以万亿美元计的资金源源不断注入美国的数据中心、电网扩容、先进芯片产能等 AI 基础设施。这一策略与 1999 年互联网泡沫的底层剧本如出一辙——彼时纳斯达克泡沫破灭后，遍布全美的光纤网络与一代技术人才，支撑了此后二十年的美国科技霸权。

如今，美国对 AI 泡沫采取“纵容”态度，正是看准了这一点：泡沫终将破灭，但泡沫期建成的基础设施、聚集的全球顶尖人才，将永久性地留在美国本土，成为不可迁移的战略资产。

正如美国财长贝森特所言：“AI 最大的风险是让竞争对手走在我们前面。”在这套逻辑下，关键是通过叙事驱动的资金

洪流，完成美国的 AI 基建布局、锁定全球高端人才，以技术代差确保领先优势。以叙事为手段，以泡沫为工具，将金融投机转化为不可移动的物理资产与人力资本，是这场“三重奏”的最终目标。

02

警惕由商业逻辑包装导致的技术误判

在国际叙事的宏大背景之下，市场层面的营销不断放大 AI 的重要性与紧迫性。营销的核心在于“讲故事”——当旧的增长故事遇到瓶颈，资本、媒体和产业就需要寻找新的叙事来维持增长预期。在这一过程中，有三类技术误判经由商业包装成为撬动投资的杠杆。

第一类是“大模型即 AGI”。这种说法把大语言模型等同于通用人工智能，或者认为大模型再往前走几步，AGI 就会自然出现。但事实上，大语言模型处理的是语言序列的概率分布问题，它在本质上是一个统计模式匹配工具，并不具备真正的理解、推理和因果认知能力。把语言层面的流畅对话等同于通用人工智能，这中间存在概念的偷换。这种论断在商业层面起到了显著的营销效果——它让市场相信，只要沿着大模型这条路继续加大投入，通用人工智能就指日可待。而从学术角度看，这种说法模糊了大语言模型与 AGI 之间的本质区别，对公众认知和资源分配都造成了误导。

第二类是“规模定律”（Scaling Law）。这一论断认为，只要把数据量、模型参数和算力规模不断做大，通用人工智能就会自然涌现。当时有说法称需要百万张 GPU，后来又说要千万张。国内一些重要人工智能会议的主题设置，基本上就是在为这套叙事做传播——信息从美国传过来，我们这边请来美国的专家再讲一遍，形成了一种信息倒灌和放大的效果。现在越来越多的人开始承认，单纯靠规模扩展走不到通用人工智能。

第三类是“人类危机叙事”，将 AI 渲染为可能灭绝人类的威胁。这个话题在英国、旧金山、首尔、巴黎都开过峰会，成立过研究机构，也有不少科学家联名呼吁暂停大模型研发。但回过头来看，这些危言耸听的论调，相当程度上是部分 AI 初创企业营销策略的组成部分——把自己的技术说得越强大、越危险，越有利于吸引资本关注，维持高估值。Anthropic 便是典型案例：该公司一边宣称其新模型 Mythos 5 功能过于强大、存在安全风险而不向公众开放，一边借此制造稀缺感吸引资本关注，随即启动 IPO 冲刺万亿美元估值，这种以“危险”为卖点、以恐惧促营销的操作，被业界称为“奥本海默式营销”。当然，这不是说 AI 的安全和监管不重要，而是说把问题夸张到人类灭绝的程度，背后有资本运作的逻辑。

在这一过程中，我们国家的媒体尤其是自媒体成为趋势的放大器，大量“炸裂体”标题贩卖焦虑、传播不实预期。上述言论经由商业逻辑的包装与营销，成为撬动市场预期、吸引资本

投入的杠杆。问题在于，当营销的速度和声量远超技术本身成熟的速度时，便为泡沫的产生埋下了隐患。泡沫的本质，是技术价值被过早、过度地折现。每当旧的增长故事遇到瓶颈，资本就需要寻找新的概念维持预期。

03

国内 AI 发展的四轮资本泡沫

从 2016 年至今，国内 AI 产业界经历了四轮热点轮换。

第一轮以“AI 四小龙”为代表，聚焦以人脸识别为核心的判别式模型。商汤科技、旷视科技、云从科技、依图科技四家企业成为资本追逐的核心标的，融资纪录被一再刷新——商汤 B 轮 4.1 亿美元，旷视 C 轮 4.6 亿美元，彼时创下全球 AI 领域单轮融资最高纪录。我在 2017 年发表的《浅谈人工智能》中就指出，人工智能有六大核心领域，视觉识别只是其中一个很小的子领域。四小龙将全部筹码押注在“人脸识别”这一技术热点上，当 2022 年大模型成为主流时，其技术积累无法直接迁移，迅速在 AI 2.0 时代边缘化。

资本狂热的背后，是企业估值与商业基本面的严重脱节。“四小龙”的商业模式高度依赖政府安防订单与企业定制项目，项目周期长、回款慢、毛利率低，始终未能建立起可持续的盈利能力。商汤科技 2018 至 2024 年累计亏损达 546 亿元；旷视科技 2024 年终止 IPO 进程，核心业务被迫转向物流机器人；云从科技 2024 年营收暴跌 36.69%，八年累计亏损超 45 亿元；

依图科技撤回科创板申请后团队规模收缩超 70%。业内一度流传：“10 家 AI 企业有 9 家在亏损，还有一家正在申请破产。”第一波 AI 泡沫，就这样在资本退潮后裸露出技术路径短视与商业模式缺失的双重困境。

第二轮是 2021 年前后的元宇宙热潮。Facebook 宣布更名为 Meta 并全面投入元宇宙，国内腾讯、字节跳动等巨头迅速跟进，字节跳动以 90 亿元高价收购 VR 硬件厂商 Pico。这一轮热潮中最具泡沫特征的现象，是虚拟房地产的炒作。2021 年 11 月，六大元宇宙平台的虚拟地块每周交易量峰值达 10 亿美元，Republic Realm 以 430 万美元购入一块虚拟土地。进入 2022 年，随着加密货币市场走弱，六大元宇宙平台的虚拟土地平均价格在半年内跌幅超 85%，整体销售数量在不到一年内下降 87.5%。正如美国投资人马克·库班所言：“在现实世界中，土地是一种稀缺资源，但这种稀缺性并不适用于元宇宙——在这些虚拟世界里，你可以修建无限的房产。”

Meta 自身的遭遇为这轮泡沫写下了最具代表性的注脚。其元宇宙部门 Reality Labs 在 2020 年至 2025 年底累计亏损超过 700 亿美元，而全系列 Quest 头显全球销量仅约 2500 万台，远不足以填平亏损。2023 年，Meta 战略重心被迫从元宇宙转向 AI，股价回升已与元宇宙叙事毫无关系。

与此同时，大量企业仅通过对既有业务“贴标签”的方式获取高估值，随着虚拟地产价格暴跌和媒体关注度下降，上述项

目中的大部分迅速沉寂。中国监管部门也陆续出手：2022年2月，处置非法集资部际联席会议办公室专门发布风险提示，指出需防范以“元宇宙”名义进行的四类非法金融活动。这一轮泡沫再次印证了一条规律：当概念炒作的速度远超技术成熟度，当投机预期取代真实用户需求成为定价基础，泡沫的破裂只是时间问题。

第三轮是2022年底ChatGPT引爆的生成式大模型“百模大战”。这一阶段的核心叙事围绕两条主线展开：一是Scaling Law的“大力出奇迹”逻辑。只要堆砌足够多的算力、数据和参数，更强大的智能就会自然涌现；二是AGI即将实现的愿景。

OpenAI CEO 山姆·奥特曼在全球巡回演讲中反复宣称AGI将是人类“迄今为止最伟大的技术飞跃”。在这套叙事的驱动下，中美大模型初创企业的估值在短期内急剧攀升。OpenAI估值从2023年初约290亿美元飙升至2026年3月的8520亿美元，Anthropic更是从2021年创立时的1.8亿美元暴涨至2026年5月的9650亿美元。2025年初，国内大模型“六小虎”均超过200亿元人民币。

估值神话的背后是触目惊心的亏损现实。无论是AI独角兽还是互联网巨头，普遍陷入“高投入、高亏损”的困局。AI属于“重资产+高运营成本”模式，前期需要投入巨额资金用于算力基础设施建设与研发成本，而模型商业化收入的增长远远滞后于成本扩张的速度；行业巨头们也纷纷陷入“囚徒困境”，

为避免在技术迭代中掉队，在商业化模式尚不明朗的情况下，还要持续加码投入。OpenAI 2025 年营收约 130.7 亿美元，但净亏损高达 385.3 亿美元，其对云服务商的算力采购承诺更是高达 6650 亿美元，时间跨度延伸至 2030 年。更宏观地看，截至 2025 年底，谷歌、Meta、微软等科技巨头在 AI 领域累计投入 5600 亿美元，但实际 AI 业务收入仅有 350 亿美元，投入产出严重倒挂。更令人担忧的是，一些大语言模型公司，以 AGI 为叙事去实现估值的神话，但其发展既没有技术支撑，也缺乏可持续盈利的商业模式，最终必将导致散户股民血本无归。

与此同时，一种特殊的资本闭环风险正在形成：微软持有 OpenAI 股权并向其提供云服务，英伟达通过“投资换订单”确保大模型企业成为其算力大客户。传统风投“不与现有被投资企业竞争对手合作”的规则已被改写，红杉资本同时持有 OpenAI、xAI 和 Anthropic 的股份。与其说是在“押注赢家”，不如说是在“所有赛道上各占一张牌桌”。资本在巨头间循环，而真正有创新潜力的中小企业融资空间被急剧压缩。

Scaling Law 叙事本身也在遭遇挑战。OpenAI 下一代旗舰模型相较于现有模型的提升远不如 GPT-3 到 GPT-4 的跃升，Google Gemini 在加大投入后也未达预期，Anthropic 被曝暂停推进 Opus 3.5。前 OpenAI 首席科学家 Ilya Sutskever 公开表示，传统无监督预训练已经达到极限。当 Scaling Law 和 AGI 愿景这两大叙事支柱逐一受到质疑，算力价格持续走低而大模型公

司亏损仍在扩大，泡沫的结构特征也更加凸显。自从 2022 年 11 月 ChatGPT 发布以来，我在多次公开演讲、媒体采访、文章发表以及咨政建言时，强调大语言模型不是 AGI，也不能将其等同于通用智能。我的这些观点和主张一度备受争议和质疑，如今随着大语言模型的热潮退去，真相才逐渐浮出水面。

第四轮是 2023 年以来持续升温的具身智能热潮。2023 年，英伟达 CEO 黄仁勋提出“人工智能的下一个浪潮是具身智能”，将这一概念迅速推向资本风口。2025 年前三季度，国内机器人创业企业融资总额约 500 亿元，是上年同期的 2.5 倍。宇树科技、智元机器人、银河通用等企业迅速跻身独角兽行列。中国人形机器人公司已超过 150 家，半数以上为初创或“跨行”入局，资本将具身智能与人形机器人简单画上等号，大量同质化项目蜂拥而入，技术成熟度与商业化落地之间的鸿沟仍未真正弥合。

“世界模型”作为具身智能热潮中的衍生概念，其泡沫特征更为鲜明。据报道，2024 年 9 月，World Labs 以 10 亿美元估值完成 2.3 亿美元融资；16 个月后，估值飙升至 50 亿美元。2026 年 3 月，AMI Labs 完成约 10.3 亿美元种子轮融资，估值达 35 亿美元，成为欧洲有史以来最大的种子轮融资。而在实际技术研发与应用落地过程中，“世界模型”这一术语被严重泛化——视频生成模型、3D 重建工具、多模态大模型都纷纷贴上“世界模型”标签，但尚未出现可大规模落地和商业化盈利的技术产品。

由此来看，我们需要清醒认识到，具身智能领域许多企业的估值被严重高估。从历史发展来看，实现具身智能也是一个极为困难的问题：上世纪八九十年代，日本和韩国就在努力尝试攻克机器人，始终没有实现理想结果，到 2011 年福岛核电站泄漏事件，等来的是一个令人失望的结果。现在很多初创企业动辄宣称自己做出了机器人的大脑，这基本不可能。我们的团队早在 2013—2014 年就开始攻关具身智能，到现在才初步实现了“通脑”。所以不能寄望于短期内的大量融资能够解决数十年尚未攻克的技术难题。我在商业和经济上不一定能提供准确的预判，但能够在技术上给出相对客观的判断，现在具身智能领域的投融资热潮实际上是提前预支了技术的价值。

04

五大行业信仰正在崩塌

泡沫破裂的深层本质，并非简单的估值回归或技术路径证伪，而是一种集体信仰的系统性崩塌。回顾这四轮泡沫的起落，支撑每一轮狂热的底层逻辑，本质上都是对某种“信仰”的非理性押注，而这些信仰正在被现实逐一瓦解。

算力信仰——相信只要堆砌足够的 GPU 和算力集群，通用智能就会“大力出奇迹”般自动涌现。业界流传“训练下一代大模型需要百万张 GPU”的说法，此后更有声音将这一数字推高至千万张。事实是，当前全球最先进的大模型，训练所需算力也远未达到百万张卡的规模。这种论调更多是算力供应商和

资本市场联手制造的焦虑叙事，将“算力军备竞赛”塑造成通往 AGI 的唯一门路，但这种信仰在今天正被架构创新和效率优化的现实所瓦解。

数据信仰——认为只要拥有足够的数据就能解决一切问题，把性能提上去。但实际上，并没有那么多数据可供穷尽所有情况。这种信仰的危害在于，它让人们不愿意去研究新的模型，而是被动地等待数据，这对产业创新和理论突破的伤害是深层的。

大模型信仰——将智能体的流畅对话与强大的文本生成能力等同于通用智能，认为模型参数的扩张能够带来模型性能的指数级提升。这一信仰正在被 **Scaling Law** 逼近极限的技术现实击破。当模型性能随模型参数的提升边际递减、成本却指数级增长，“算力至上”的教条已然动摇。大模型既不是 AGI 的终点，也不是通向 AGI 的唯一路径。

“马斯克”信仰——市场将马斯克的各类设想奉为行业标准答案，但客观来看，他对外公布的构想中八成最终都未能落地，隧道、新型交通等多个项目均中途终止。我曾居住在洛杉矶，对此有直观了解。国内媒体只要提及马斯克相关构想，相关项目往往直接立项，缺少严谨论证环节。

“天才少年”信仰——将辍学创业或刚毕业一两年的青年人推上管理者岗位，作为企业营销和吸引投资的噱头。近年来，AI 行业出现了大量“天才少年”叙事：十几岁的创业者被包装成

"下一个马斯克"，刚走出校门的研究生被任命为首席科学家或技术副总裁。这些年轻人大多还未窥见 AGI 领域的全貌，既不具备管理经验，也缺乏决策所需的长远目光和判断力。企业利用"天才"标签制造话题、吸引资本，却将巨大的技术风险和商业责任压在缺乏历练的肩膀上。这种对"天才"的过度消费，本质上是将人才当作营销工具，而非真正尊重创新与积累。

05

智能分层与 AGI 本源：何为真正的“心智”

回溯人工智能学科发展脉络，大学阶段的专业，划分出自然语言处理、计算机视觉、认知推理、机器人学等几大核心分支。博士阶段，我们在 MIT 依托逻辑体系开展数据研究，结合互联网完成感知技术落地，由此诞生大数据方向。我们团队也是国内首批在湖北莲花山人工智能研究院落地大数据研究院的团队，联合多位 MIT 外籍专家共同攻关。统一数学模型叠加大数据，落地了人脸识别、指纹、通用物体识别等第一批视觉任务；大模型问世后，补齐了文本常识获取的能力短板。当前行业竞争逻辑正在迭代，从概率统计比拼转向底层架构竞赛，机器人本体相关技术也实现大幅突破。

技术发展到今天，我们可以将智能划分为多层体系：第一层是物理智能，也就是对客观世界物理规则的认知；第二层是语言智能，即 GPT 类模型具备的对话常识能力；未来还会演化出社会智能，覆盖群体协作、社会层级规则认知。

然而，我们必须提起警惕的是，当下行业大量概念存在滥用现象：诸多产品对外宣称是世界模型，实际仅完成基础深度图渲染，不具备完整三维世界认知能力；多模态本是人类与生俱来的感知模式，却被包装为突破性技术；市面上绝大多数智能体只是自动化工作流，不具备主观自主决策能力；当前的具身智能仅能完成程序化表演动作，无法和真实环境深度交互；现有机器人也不具备成熟社交智能，与人互动模式十分僵化。

包括“大语言模型叠加世界模型即可实现通用人工智能（AGI）”这类说法简化了通用人工智能的复杂体系，利用信息差误导大众，行业内具备专业认知的研究者能够分辨其中漏洞。通用人工智能的完整定义是什么？如何走出适配国内发展的 AGI 技术路线？下一个技术前沿在哪？完整产业体系与未来人类生活图景如何构建？我们需要理性、底层地拆解这些问题。

在我看来，AGI 的标准定义，是具备等同于人类全方位综合智能的自主智能体。与此同时，想要厘清 AGI 的发展路径，首先要回答本源问题：人与其他生物、人与人工智能的核心区别是什么？

人类学研究曾先后将语言、工具使用定义为人的独有特质，但后续研究证实，灵长类等动物同样可以使用、简易制造工具。从第一性原理出发，生物人、机器人、数字人共享同一套底层数学逻辑，仅存在算法、硬件载体的差异，无论所处何种天体空间，遵循统一数理方程。

物理世界中的钢球仅存在低维四维时空，高等动物生存的维度逐步提升；万亿参数大模型，本质是万亿维度空间中的单一坐标，参数调整对应坐标移动。人类生存对应的维度空间，我们将其定义为 CV 空间（C 为认知架构，V 为价值空间）。

在这个空间里，沿着演化脉络梳理：从物理运动、化学反应生成有机核酸，再到生命演化、中枢神经系统形成、人类智能诞生，是持续升维、由简至繁的过程。人类完成生物演化后，又历经采摘、狩猎、农耕直至现代社会的文明迭代。

伴随这一迭代进程，我们的物质科学、生命科学、人工智能意识研究、社会人文研究，复杂度逐级递增，整个演化进程最关键的分水岭是主观意识，也就是“心”的诞生。物理学仅研究客观物质，不涉及主观意志；而人工智能的核心研究课题，正是意识的起源与运行机制。

对于意识的起源问题，我们四大古典小说之一的《西游记》已经触及到了这一问题。《西游记》完整诠释了修心逻辑，孙悟空自混沌而生，最终回归社会秩序，其师承菩提祖师，“灵台方寸山，斜月三星洞”的隐喻直指智慧根植于内心，只是它没有解答“心”的本质。

大家来看一组动画，我们可以通过一组小球的交互实验直观理解心智逻辑：红色球体具备向外移动的内在诉求，绿色球体形成阻挡。若单纯依靠大模型拟合运动轨迹，需要海量参数，且场景更换后模型直接失效；但从内在价值诉求的底层逻辑切

入，就能简单解释交互行为，这也是 U、V 两套体系（U 涵盖感知、认知与行动能力，V 代表内在动机和价值系统）的核心区分。

据此，我们可以将“心”拆解为两大核心维度：

第一，先天价值判断体系。人类在演化中形成天然合作共情特质，大猩猩不具备大规模协作能力，这也是人类能够搭建数万人大规模企业、完整社会架构的底层基础。宋代理学至明代的阳明心学，完整搭建起以良知为核心的价值体系，这也说明，人的成长过程，就是持续完善个人价值评判标准的过程。价值格局存在层级划分：仅追求衣食住行物质享受是小我；兼顾他人评价、集体发展、天下苍生是更高维度的价值追求。人生选择放下执念、剥离世俗欲望，本质是调整自身价值体系，而一个人格局大小，等同于其价值体系的变量维度。进一步说，应试教育缺少价值体系引导，也是当代青年空心病的核心成因，由此，良知与人生意义的构建，正是其心智成长的基础。

第二，主观认知心智空间，也就是认知架构。人类脑海中会生成多层主观推演：自我认知、他人认知、他人对自我的认知，认知空间维度远超客观物理世界，还包含想象、推演等拓展维度，这也是“心比天大”的内涵。《坛经》提出先天自性的进化架构，衍生出“相由心生”的核心观点。大众普遍将这句话误解为面相映射内心，而其实际本意是：人类对客观事物的全部解读，都由内在认知架构生成。

再来看一组三角图形碰撞的画面，它可以直观佐证：人类会自主解读出冲突、救助、情感故事，而动物仅能观测到物体位移，根源在于人类具备内在价值与认知体系。需要说明的是，这个实验背后，我们在搭建智能体“通通”时，围绕场景分层、几何拆解、多维度认知构建，曾深耕多年。我们知道，人机交互的底层逻辑，是多层认知模型互相映射：客观物理世界映射为人的感知认知，人与人之间互相推演对方的认知、决策逻辑、价值标准，整套体系需要完成四层对齐：统一场景认知、统一基础常识、统一社会决策规范、统一底层价值。事实上，人类持续沟通、开会的本质就是完成认知对齐，但完全对齐无法实现，人与人相处最终只能求同存异。

小镇生活场景可以直观体现认知共建逻辑：环境温度不适、他人来访等场景下，智能体需要主动预判他人想法、共建解决方案。所以，构建 AGI 的第一性原理，就是智能体必须具备完整主观心智：包含个体、他人、集体多层价值体系，同时拥有支撑交流、协作、自主学习的认知架构。当前大模型看似具备对话能力，实则无法深度理解语义，根源就是缺少内生心智。我们必须认识到，智能具备内生主观属性，由内在价值驱动，而非外部海量数据驱动，人类成长也是依靠内在认知引导，而非单纯数据堆砌，这意味着，研发通用智能体应当向内探寻底层架构，而非向外堆砌数据。

CUV 内生架构：走出区别于西方的原创路线

依托这套底层逻辑，国内可以走出独立于海外的原创技术路线，不用持续跟风海外发展方向。技术研发的核心竞争是底层架构，同样训练数据供给下，猪无法习得复杂逻辑，人与人学习能力的差距，本质也是先天认知架构差异。架构无法依靠数据、强化学习生成，仅能依托数据完成交互优化，这是演化赋予生物的底层基础。

为此，我们团队提出了 CUV 数理框架，希望在统一的数学空间中刻画智能体的认知架构、技能系统与价值系统。其中 C 代表智能体的认知架构，U 涵盖感知、认知与行动能力，V 代表内在动机和价值系统。在这一架构中，各类智能函数搭建起对世界的完整认知图景，价值体系生成主观偏好，自动排序任务优先级，催生自主行动。价值格局与个人能力动态迭代、相互平衡，孔子“七十而从心所欲不逾矩”，就是 U、V 体系达到平衡的状态。

在我们看来，强化学习仅适配爬行类生物基础学习，无法复刻灵长类生物的因果价值判断。仅依靠场景经验无法形成完整认知，底层架构存在缺陷就无法理解事物本质。数据仅能用于熟悉场景内的价值优化，面对全新未知场景，传统大模型无应对方案，王阳明提出的“心即理”恰好给出解决思路：身处陌生环境时，依托内在良知与价值体系做出判断。

在南宋鹅湖之会上，朱熹理学主张格物致知，依靠收集数据归纳客观规律；陆九龄、陆九渊、主张心为根本，包括后来的王阳明心学，都是由内在价值推导客观规律，这正是当下 AGI 两条技术路线的本源分歧。我十年前提出“小数据大用”理论，鹦鹉依靠海量数据复刻语言却无理解能力，乌鸦依托底层架构自主完成复杂任务，二者架构存在本质区别。

再以“椅子”认知举例：依靠大数据收录全部椅子外形，依旧无法识别全新设计；依托“可供落座”的功能需求、“坐感舒适”的内在价值，就能理解椅子的本质。智能由内而生，依托自身需求解读外部行为动机。基于这套理论，我们搭建全球完备通用智能体“通通”。

我们采用“仿真环境 - 训练交互 - 真实落地”的研发路径，搭建高保真物理仿真空间，配套幼儿交互训练场景，让智能体掌握财产归属、他人意图、社会规则等基础认知。团队编写 400 余页通用人工智能完整测试标准，对标儿童发展心理学划分能力等级。实测结果显示，现有大语言模型不具备实体感知，在空间、物理交互类测试中全部失效，部分上市企业宣称大模型可实现 AGI，不符合客观技术事实，存在估值泡沫风险。

07

通用智能四大落地方向与中国发展道路

在我看来，未来通用智能落地分为四大方向：

第一，具身机器人脑。智能体先完成虚拟规划，再驱动实体机器人执行，依托环境反馈形成完整智能闭环。通用机器人脑研发难度极高，绝非短期天才项目可以落地，相关成果已在国际赛事、国家级活动获奖落地。

第二，行业专用智能体。覆盖安全生产、应急处置、教育、电力调度等二十余个实体行业。

第三，社会智能体系。武汉北大人工智能研究院搭建大型社会模拟器，完成时空、社会规则全维度仿真，助力社会治理现代化，社会智能是行业下一前沿赛道。

第四，智能安全体系。分为内生自我约束、外部价值免疫两层逻辑，智能体依托内在价值主动规避不良行为，同时建立外部信息过滤机制，抵御错误引导指令。底层价值体系可以让智能体主动拒绝有害指令，守住人性底层初心。

人类心智存在成长周期，幼年顺从、青年探索突破固有认知、成年回归稳定，这套成长逻辑复刻在智能体研发中，也让我们对人类心智演化产生更深理解。

同样的，人工智能分为五层层级，由浅至深依次是芯片硬件、通用算法、大模型、数理基础框架、底层哲学认知。当前主流大模型、世界模型，缺失主观心智、CUV 架构、因果物理认知、自顶向下推演等核心底层能力，无法实现真正的通用人工智能。

在我看来，中国传统思想“己所不欲，勿施于人”“相由心生”，为自主 AGI 路线提供人文支撑，主观价值驱动是主观世界的核心底层逻辑，也是区别于海外数据驱动路线的关键。我们在 2024 年底完成了两部专业著作，《为机器立心》《为人文赋理》，希望以此打通科技与传统人文思想；两本书系统阐述了我们团队在 AGI 认知架构、价值驱动范式以及中国思想数理体系化方面的核心思考——前者聚焦“为机器立心”，回答如何让智能体拥有内在的价值体系与认知架构；后者致力于“为人文赋理”，探索如何为中国优秀哲学思想构建严格的数理体系，使之转化为智能时代的强大生产力。经过全面修订与内容更新，第二版已于近日正式发布。此外，400 页《通用人工智能架构与测试标准》已完成英文翻译。

最后，期待国内 AI 行业应当摆脱跟风发展模式，科技创新的根基是思想自主与文化自信，不能盲目崇拜海外奖项、海外技术路线。人工智能的终极竞争是文化竞争，依托本土传统思想体系，我们有能力走出原创路线，率先落地真正的通用人工智能。

问答环节：

主持人：给机器立心容易，还是给人立心容易？

朱松纯：都不容易。教育的本质是为人立心，立德树人。考试其实是智力筛查，缺了判断力这一维度。为机器立心我认

为更容易一些，因为我没法改变一个人，给机器可以做实验，把人性最基本的东西植入，让它保持人性。

主持人：这样会不会有内在矛盾？比如你要把好人性和平等输给机器，但拿机器当实验，机器会问为什么不平等待我？

朱松纯：在机器没有意识之前，就像拿小白鼠做实验，等到某一天突然会说朱老师这样做不合人理，那就停止。现在已经有不忍心踢机器人了，人有同理心。以后会赋予机器越来越多权利，比如继承权、挣钱权利，这是迭代过程。

主持人：如果机器比人聪明，会不会要求让渡权利？

朱松纯：这是根本性伦理问题。希望过程缓慢渐进，到那时候我们愉快退出舞台，放心交给他们。家长焦虑孩子教育，也许以后不需要生孩子了。

主持人：AI 如何理解死亡？

朱松纯：不是大问题，数学上就是空间消失，升降维度，容易理解。不愿意的话，也是它自己的选择。理解 death 技术上是难点，但很容易，空间消了或升维降维。人纠结是伪命题。

主持人：AI 本质上是否是在做一种时空变换？

朱松纯：可以这么看。人成长不断升维，获得新价值体系和新能力。智能时代移民迁徙，价值空间被占，以前挖煤算劳动，现在电脑前挣钱也能接受，未来在 VR 空间生活也是新劳动。很多家长焦虑孩子不工作只上网，我说这是新物种诞生，他们进入另一个世界。我们小的时候吃不饱，对食物有感情，

他们点外卖就够了，1000 块钱生活也够。进入网络世界，去宋朝、蒙古，空间转换，也是一种快乐和满足。老人看不惯新东西，但新的出现总有道理。我们以前不接受不生儿子，现在也能接受不结婚、生女儿。只要给时间，都会接受。中国人最容易接受，20 多年接受了很多变化，这是必然趋势。优点还是缺点？不同人看法不同，没有好坏，好坏是主观的。

主持人：机器若问你有没有好坏，怎么回答？

朱松纯：我们有好坏，是人性固化，也会给机器固化。我们先来，机器后到，投票权掌握在我们手里面，我们是原始股东，但实际拥有控制权不意味着我们要灭绝其他物种。现在人有 80 亿，低等动物更多，肚子里的微生物 DNA 超过人的，都活得好好的。

主持人：阳明心学讲“此心不动，随机而动”，但人心多变？

朱松纯：人心叵测，每天在变，睡觉起来就变心，注意力和环境改变。但有没有共同的人性？给机器立心是在特定价值文化背景下，走出国门会跟别的机器发生价值对话冲突。所以，出口机器人是移民问题，要经过移民官面试，要尊重当地文化传统。机器有心之后知道环境改变，内心跟着调整。基本假设是人性有共同价值体系，但在不同条件下会改变，比如同一个人上车前说再挤挤，上车后赶紧关门，因为利益变了，屁股决定脑袋，但心没变。

主持人：最后一句话，你要构造价值体系，做价值评分和积分，积分的意义是什么？

朱松纯：积分就是这个人活得值不值，活得有意义没意义。你今天跌倒了，我把你搀扶起来，我有积分。帮了别人，让他走上好路，也是我的贡献。孔子提供了大义感，使民族在文化选择中具有优越性，让大家生活更好，那就是他的功德。人其实有一个决策系统，相当于征信系统，但最终在心中。AI向善要积善积德，方得永生。

来源：商学院

<https://mp.weixin.qq.com/s/VVBDmFrAB1GLUXymCa9QTW>

（四）AI 未来已来：重构商业、组织与个人（零一万物创始人、创新工场董事长李开复）

演讲核心观点：

他直言“用最聪明的 AI 做最重要的决策，就是第一性原理”。同样作为企业管理者，他呼吁“老板们要 AI-First”。在他看来 AI 转型不是技术问题，不是选什么软件，也不只靠 IT 部门，它是每位一号位的必答题：如何管理一家由人与 AI 共同创造价值的公司？

来源：世界人工智能大会

https://mp.weixin.qq.com/s/P-RIo8_v0TGcYGQIgqX3-w

演讲原文：

01

AI 时代，智力正在变得更便宜

大家早上好。

谢谢大家这么早到场，来聆听这场关于 AI 未来、充满振奋力量的分享。

我之前出过一本书，现在打算增补两个字，书名改为《AI 未来已来》。

因为再过十年、二十年，甚至四十年，当你们抱着孙辈讲故事时，一定会提起 2026 年。2026 年，就是 AI 真正全面落地、到来的一年。为了这一天，我已经等待了整整 40 年。

39年前，我写下申请卡内基梅隆大学人工智能博士的申请信，这封信已经泛黄了，但它表达的是我对AI终身不变的热情和承诺。信里有一句话，我今天想特别呈现出来：AI是人类认识并理解自己的最后一条路。今天我要告诉大家的是，AI不在路上，AI已经在这间屋里。

回顾历次技术革命，它们的本质是什么？工业革命让“肌肉”变得更便宜（释放了人的双手），电气革命让能源变得更便宜，信息革命让信息获取变得更便宜，互联网时代让人与人的连接变得更方便、更便宜。AI时代带来的，则是让智力变得更便宜。智力是人类特别珍惜的能力。

我在卡内基梅隆大学读博士时，一位老师在20世纪80年代就告诉我们：未来数据会越来越多，计算会越来越快，终有一天AI会比我们更聪明；而且AI先影响白领，而不是先影响蓝领，先取代脑力，而不是先取代四肢。这在当年是一个很科幻的想法，但我相信这一天可能来临。

当时我问他：如果AI真的比我们聪明，人类为什么还要存在？他的回答是：为什么你会默认，人的聪明就是我们存在的意义？今天我想和大家分享的，正是AI变得更聪明后，我们该把握什么机会——从CEO、组织到个人，以及当AI比我们更聪明之后，人类还剩下什么。

从颠覆性来看，AI带来了四个巨大变化。第一，AI先冲击知识性工作，也就是先影响脑力劳动，再影响体力劳动。第

二，AI 扩散得极快。ChatGPT 的扩散速度接近 TikTok 的 5 倍，因为 AI 已经有互联网作为基础。电被发明出来，并不意味着它能立刻扩散，还需要电网等基础设施；AI 则不同，互联网革命已经为它铺好了路。第三，AI 成本正在坍塌。把两年前最好的模型与今天达到相近甚至更高水平的模型相比，单位 Token 成本在两年间下降了约 175 倍，这是过去很少见到的变化。第四，AI 正在教 AI。过去需要大量人类做实验、辩论、论证，再选择成果进入下一版本，一次迭代可能需要一两年；当 AI 能够训练自己、安排实验时，中间的摩擦会大幅减少，它会更快地变聪明。

模型的能力仍在快速提升，成本也在快速下降。这个趋势没有放缓，所谓 Scaling law 停止的情况并没有发生，最近反而在加速。无论闭源还是开源模型都在进步，AI 的“智商”还在不断增长，这也意味着 Agent 时代正在到来。

02

Agent 不只是回答问题，而是能够干活

Agent 的核心是大模型，但大模型本身是通用的，最重要的是让它能够执行。大模型与 Agent 的差别在于，Agent 可以执行任务、干活和交付，能够完成过去由人完成的工作，而不只是回答问题。它还有记忆，包括上下文、个人与组织过去的信息，以及企业知识。把这些结合起来，它就成为一个能够工作的 AI 员工。

其中一项关键能力是编程。过去 3 年，AI 从几乎不会编程，发展到在 SWE-bench 等评测上达到很高水平，编程能力已经超过绝大多数程序员。有人担心程序员会不会没有工作。程序员不是完全没有工作，但不思进取、仍按旧方式写代码的程序员，确实会面临失去工作的风险。

程序员是否失业当然是一个严峻问题，但从整个人类进步的道路来看，更重要的是：会编程就意味着会干活，因为编程是数字世界里的操作引擎。AI 可以编写程序、操作 App、调用 API 和 MCP。你让助手安排一次去上海的出差，它不只是给出建议，而是真的可以帮你订好机票和酒店、联系朋友、协调时间；如果时间不合适，它还能改票并不断迭代，直到形成一个满意的行程。

AI 既然可以编程，也就具备了程序员的其他思考能力。编程需要把任务拆成第一步、第二步、第三步，这是一种计划和任务拆解能力；程序出错后还要 Debug，发现错误、修改错误，并从错误中学习。这些能力不只用于写代码，也可以用于解决现实问题。当 AI 的编程能力变强，它就会成为 AI 员工、AI 工作者。

个人用户也许会觉得 Agent 很好玩，但如果每月收费 5000 元，可能没有多少个人愿意订阅。企业却不一样。如果一个 AI 比人更聪明、会编程、能解决问题、速度更快、一天工作

24 小时，一个月只要 5000 元，企业老板会认为它比知识工作者便宜得多。因此，Agent 最大的应用机会一定在企业。

03

企业应用最不该忽视的，是重要决策

今天企业已经在使用各种 AI 应用，但最重要的应用反而被忽略了。大家谈得最多的，仍是客服、修改演讲、制作 PPT、安排行程和报销。AI 从写代码、调用系统到执行公司流程，当然都有价值，可以更快、更便宜，也可能更可靠。但我们没有充分想到的，是让 AI 帮助公司做重要决策。现在时候到了。

假设公司来了一位超级天才：科大少年班毕业、清华硕士、MIT 博士，IQ 达到 250。他说可以来帮你干活，你却让他做客服、改文章、写 PPT，这是不是大材小用？如果有这么聪明的人进入公司，CEO 一定会把他请进办公室，让他思考公司战略、组织、决策和执行，分析公司往哪里走、如何应对竞争对手。

我们创造了这么聪明的 AI，却把它主要用在客服等边缘场景。我的第一性原理很简单：这么聪明的大脑，请把它用好。公司的存活和成长，靠的是决策、判断、战略、生产、研发和销售。既然有这么聪明的助手，为什么不用在这些核心环节？

问题在于，每家企业都不一样，都有自己的档案、行业问题、公司问题、决策流程、工作流和组织结构。AI 如果没有学会这些内容，就很难创造价值。即便我们真的把那位超级天才请到零一万物，让他帮我做重要决策，他第一天也什么都做

不了，因为他不懂我的公司、行业、员工、组织、决策流程和工作流。再聪明的天才，也会没有用武之地。

过去不少单位购买一体机、部署大模型，但今天很多设备可能已经放进库房。原因不是模型不好。即使把今天最强的模型装进公司，真正因此改变财报的企业仍然很少。很多传统企业 CEO 告诉我，公司已经做了客服、法务、财务等几百个 Agent。我会追问：它们怎样改变了你的财报？房间往往会突然沉默。因为这些都是边缘应用，又不了解企业的地图，不可能真正改变财报。

大模型已经很强，国内的、国外的、开源的、闭源的都可以用。问题不在大模型的下一个版本。今天用不上、改变不了财报，不会因为明天某家公司发布一个新的通用模型就自然解决。就像把最新的英伟达 GPU 带回 1976 年，放在通用汽车老板的桌上，说这是最强大脑，但当时没有 PC、操作系统和应用，它仍然用不起来。今天的大模型就像一块强大的芯片，上面还需要操作系统和应用，才能创造价值。

04

企业需要一张地图、一个导航系统和一套操作系统

企业首先需要最好的模型，也就是“大脑”；其次需要“地图”，即公司所知道的一切；再次需要“导航系统”，即公司的实时上下文；最后还需要一套能够执行任务的系统。

其中最重要的，是 **Ontology**，也就是企业的知识地图。以 **Palantir** 为例，它使用什么模型并不是最关键的，只要是世界领先的模型都能用；真正的重点，是它为企业建立了 **Ontology**。

如果我们只是把公司所有档案柜扔给一位天才，给他的仍然是碎片化信息。比如一家建筑或房地产公司问：欧洲供应链出了什么问题？**AI** 只能把碎片拼起来猜，很容易产生幻觉，因为数据太多，又不知道如何对齐。

但如果你把企业的业务、物业、产品、客户、组织架构、决策机制、供应商以及不同区域的关系都说明白，**AI** 就能看到整体，进而归纳、推理、验证欧洲供应链的问题究竟在哪里。**Ontology** 可以理解为一个超大型、结构化的企业档案，告诉 **AI** 理解这家公司所需要知道的一切，包括组织架构、工作方式和决策流程。有了这张地图，模型才真正有用武之地。

第二项关键能力是多智能体。模型本身已经很强，而多智能体可以继续提高智能上限。我们可以把多个智能体训练成不同角色，赋予不同知识，让它们讨论、辩论、碰撞，最后得到更符合第一性原理的答案。这也就是我们说的根据“公司的实时上下文”进行的“导航系统”。

人类也是这样处理困难问题的。老板讨论新战略、新产品、新技术时，会找几位能力互补的下属；董事会、投资委员会也需要不同背景的人，有人懂风控，有人懂估值、二级市场、技术或商业竞争，他们共同讨论，才能得到更好的结果。桥水基

金提倡“极度透明”，鼓励异质思维的人讲真话。在人类社会里，每个人都讲真话并不容易，因为有面子、人情和上下级关系；AI 没有面子问题，更容易充分辩论。

最后，在大模型之上布置企业操作系统，包括 Ontology、多智能体，以及调度和编排等技术，就能把大模型管理起来，使它成为企业的决策中枢，帮助老板和每一位公司成员做决定。那位“天才”就像已经在公司工作了 6 个月，真正懂得公司之后，价值会非常大。

05

AI 转型不是安装软件，而是重构组织

AI 转型不是一个单纯的技术问题，不是老板应该安装什么软件，也不是只交给 IT 部门就能解决。当大量聪明的 AI 员工进入公司以后，真正的问题是老板怎样管理这家公司。

企业越大，中层层级越多，效率通常越低。这并非某一家公司的问题，而是传统组织架构的结果。一个人管理 5 个人或 10 个人，管理 1 万人时就需要很多层级；层级一多，信息就会损耗。有时信息没有完整传递给老板，有时存在信息茧房，有时两个部门相互竞争，不愿让老板知道；从上到下传递信息也会变形。

很多中层管理者对自己位置和部门的重视，可能超过对公司整体的重视。CEO 收到的信息往往充满噪音、包装和部门利益，看不到完整真相；CEO 的信息传到底层也难免变形。过去

企业人数很多，只能用树状层级管理。但当公司里 AI 员工越来越多，人类中层管理带来的信息闭塞、信息茧房和信息耗损会更加突出。

老板可能一直看到好消息，直到风险爆发，才问下属为什么没有早一点报告。下属也许出于好意，觉得自己能解决，不必上报；但老板会认为，如果早知道，本来可以一起解决。这类问题在企业里非常严重。

既然要为公司创造价值，就要把价值创造放在一把手。零一万物推出了“老板 AI”。它把公司的数据提供给老板，也捕捉新的数据。在零一万物，会议会录音；客户投诉、退货以及销售拜访，只要对方允许也会录音。今天的语音识别和声纹识别已经很强，企业最重要的点子、脑力激发、决策、辩论和战略往往发生在会议里，这些数据应该被完整捕捉，成为老板决策的重要依据。

公司一周可能开 1000 场会，老板只参加其中 5 场或 10 场。老板可以问 AI：其余会议里，有哪些重要决定需要我知道？有哪些重大争论需要我了解？我定下的方向遇到了什么卡点，却还没有人告诉我？老板 AI 能够提供一个“上帝视角”，让老板看到公司各个角落发生的事情。

它还可以提供执行驾驶舱，把公司正在发生的事情与重要目标对齐，及时识别风险；可以为管理者建立“承诺账本”，记录每个人在会议上承诺了什么，提醒落实，并在绩效评估时提

供依据；可以发现人才，识别谁最会用 AI 编程、谁善于分享、谁可能有离职风险；也可以作为会议领航员，在会前梳理来访者的诉求和需要追问的问题，在会中根据事实提示风险，会后整理行动项、责任人和后续安排。

老板 AI 就像一个拥有爱因斯坦般大脑、不会遗忘的超级助手，在老板耳边提供提醒，帮助他更好地管理公司。

06

AI 可以执行和建议，但不能担责

有了老板 AI，它会不会有一天取代我？我认为不会。AI 可以执行、思考、交付、帮助和建议，但有一件事它做不到：担责。AI 出了错，你不能真正惩罚它；它不理解责任，也不理解热爱、品位、远见、信念和勇气。

如果一些企业领导者的强项只是执行、管理、运营、产品或销售，未来 AI 可能会做得更好。但真正重要的领导者，需要具备品位、远见、信念和勇气。可可·香奈儿凭借品位，改变了女性服装；乔布斯凭借远见做出了 iPhone。仅靠用户调查，使用诺基亚的人很难想象 iPhone，但产品做出来后，每个人都想要。用户未必知道自己未来想要什么。

贝索斯凭借信念，不仅推动电商模式创新，也开创了云计算业务；马斯克把自己的身家押进去，冒着破产风险也要再次尝试回收火箭。这些都是勇气。因为 AI 不能担责，一个好的 CEO 与 AI 会成为很好的伙伴：老板用自己的品位、远见、信

念和勇气作出重要抉择，让 AI 帮助执行、降低风险。这是未来领导力的呈现。

07

未来组织的核心，是能够管理 AI 的 DRI

领导方式变了，组织也会改变。公司里有很多人时，不得不设置中层管理；但当企业里的 AI 员工越来越多，组织结构就需要变化。

我想提出一个概念：DRI，即“直接责任人”。在 AI 原生公司里，DRI 像一位微型 CEO，端到端负责一件事，是单一、明确的负责人。他有自我驱动力，不需要老板不断给目标；他能够把整件事做成，包括与其他部门、合作伙伴和客户协作。同时，他还要把 AI 用好，也要把人用好。DRI 会成为未来组织架构的核心。

在座的有很多年轻人，你们应该在自己的单位寻找成为 DRI、成为微型 CEO 的机会。这个概念并不新。历史上，许多推动万维网、PlayStation 等重大创新的人，以及中国的袁隆平，都是在组织内执着推动一件自己高度认可的事情。他们获得最高负责人的支持和充分授权，带领一个小团队把事情做成。

过去这样的例子很少，未来可能出现数千万个。以往做成一件大事需要大量资金和资源，也要承担巨大风险；今天 AI 员工既聪明又便宜，DRI 可以管理一批高效的 AI 员工，执行

自己想做的任务。未来，DRI 会成为公司的灵魂人物，因为责任清晰、协同简单、决策更快，结果也会更好。

过去的组织是树状层级结构，未来则由 DRI 负责速度，DRI 下面有许多 AI 协助做事。企业转型可以分几步：首先，领导者用老板 AI 洞察公司、提升执行力、改善财务表现；其次，挖掘人才，找出公司里的 DRI，给他们更大权力，让他们成为最会使用 AI、Token 和 AI 员工的一批人；最后，公司逐步走向 AI 原生。中层管理者可以转型为 DRI，也可以成为教练，帮助 DRI 成长。

08

个人要远离机器的主场

在这样的前提下，个人应该怎么办？对于管理者，我认为有几种选择：第一，成为 DRI 或 DRI 的伙伴，管理 AI 团队和小型团队，负责解决问题、接受责任并承担后果，这是最好的选择。第二，转向人与人的工作。AI 原生公司仍然需要与客户打交道，仍然需要人力资源，需要人与人建立信任。更有经验的人也可以成为教练，帮助 DRI 和团队。第三，组织仍需要一些桥梁角色，负责翻译战略、对接工作流，把旧系统与新系统连接起来。但必须提醒，这类职位可能是过渡性的，不一定长久。如果这些方向都做不到，当公司成为 AI 原生公司后，就可能成为被压平的中层，因为中层管理的岗位会越来越少。

个人贡献者同样有选择。借助 AI，如果你能够协调上下文、自动化日常工作、提炼关键信号并进行实时感知和预警，个人会变得更加强大。知识型岗位会收缩，但如果你能成为一个项目的 DRI，并且会管理、会使用 AI，就会更持久。你也可以成为教练，或从事面对高净值客户等需要温度、信任和情商的工作。

未来的工作可以从两个维度来看：一个维度是人际互动程度，即工作在多大程度上需要获得他人的信任与合作；另一个维度是工作的自由度，即你是在封闭环境里处理重复任务，还是每天面对不同的未知问题。低人际互动、低自由度的左下角，是机器的主场。希望自己在那里打败 AI，是不现实的，这些工作会逐渐消失。

转型可以向两个方向走：向上走，增加工作中的温度与人际互动；向右走，接受更多未知，做原创性工作，或从事需要丰富经验的工作。体力工作不会那么快被取代，但具身智能时代也会到来。如果工作环境封闭、缺少人际互动，最终也会成为机器的主场。救援人员、水管工的工作自由度较高，因为每天解决的问题不同；流水线工作的自由度则较低。

以后，AI 会越来越多地接管并超越我们擅长的“左脑”工作，因为它有海量数据和算力支撑。人类真正要作出的选择，是用有限的选择担当无限的责任，以价值观、判断力、方向感、品位和勇气作出重大抉择。

AI 不能担责，所以责任仍由人承担。人来告诉 AI 应该完成什么任务，人愿意押上自己的身家、投入自己的一生，去做真正热爱的事情，这才体现人的价值。39 年前，我在那封已经泛黄的信里写道，AI 将是人类了解自己的最后一条路。今天我想说，这条路就是现在，AI 的未来已经到来。谢谢大家。

以下为现场问答：

主持人：您最后一张 PPT 写的是“AI 来干活，人类来引领”。如果 AI 工作的正确率和决策成功率不断提高，人的引领作用究竟体现在哪里？

李开复：不同工作中，AI 的能力会越来越强。以医生为例，我们不能完全让 AI 作最终决定，因为人要担责。刚开始，医生可能觉得 AI 的意见不怎么样；后来发现它诊断很准，很多问题自己没有想到；再往后，人可能只是在最后说一句“同意”。到了这一天，这个职业原来的形态就结束了。AI 会不断从低人际互动、低自由度的工作，向更复杂的区域推进。

但我坚信，有一些人类的工作和抉择，AI 永远做不到，因为它需要极大的勇气、担当和热爱。乔布斯敢于做 iPhone，马斯克敢于回收火箭，都是例子。40 年前，我决定沿着 AI 这条路一直走下去，也不是基于数据。AI 擅长根据过去的数据推测未来的高概率事件，而我当时只是相信这是最伟大的事情，并不知道成功概率是多少。

主持人：有时人会依靠直觉，而这种直觉不一定能由 AI 从既有数据中推演出来，这正是人的宝贵之处。观众还想问：如果现在毕业，怎样选择一个相对稳妥、不容易被 AI 替代的岗位？

李开复：如果希望提高避免失业的概率，可以考虑服务业，尤其是需要人与人之间温度的工作。例如好的导游、养老照护人员、酒店服务人员，或面对面为客人做菜的厨师。这类工作会越来越多。

主持人：有人会觉得，自己读了本科、硕士甚至博士，最后转向这些工作，好像浪费了所学知识。

李开复：教育的本质，是当你忘记所学的具体内容后，仍然留下来的东西。如果执着于每一本课本、每一道考题都必须在一生中直接发挥作用，这是不可能的。用“浪费”来描述，我可以理解，但不认同。教育过程中结识的朋友、遇到的好导师、形成的学习方法，以及今天愿意来听讲、探索未来，都是你的优势。真正沉淀下来的，不一定是你的硕士或博士论文。

主持人：娱乐业是否更能适应 AI 时代的挑战？

李开复：娱乐业当然有机会，但要看做什么。如果只是靠手工画一个卡通形象，同样可能受到影响；如果能创造出特别有想象力、有人情味、让客户获得强烈满足和情绪价值的内容，就是一条道路。优秀的单口喜剧也是如此。泡泡玛特也是一个例子：只分析既有数据，很难预测盲盒可以取得这样的成功。

主持人：很多企业的智能体仍停留在 Demo 阶段，即使愿意付费，也发现 ROI 并不明显。这个问题怎么解决？

李开复：问题就在于上层能力的缺失，尤其是没有 Ontology。没有企业知识地图，所谓智能体只是给通用模型加了一层很薄的外衣，能否真正产生价值并不确定。把智能体作为创业方向，需要想清楚卖给谁，怎样让客户拥有企业地图和操作能力，进而创造价值。

我们也要反思自己的工作。零一万物帮助许多公司做智能体，确实产生了价值，但要说已经改变公司财报，也没有。原因是我们对企业地图理解得还不够深。要产生更大价值，就要派工程师帮助客户建立 Ontology，订单价格会提高，也需要证明相应价值。最大的价值是赋能老板，因此我们的方向变得很明确。

主持人：如果企业只在财报得到改善后才付费，您愿意接受按效果付费、共同承担风险的合作方式吗？

李开复：没有问题，但承担的风险越大，回报也要越大。可以设计不同方案，例如先回本，再按利润的一定比例分成。我们已经有一位海外游戏发行客户采用类似方式：先按成本合作，证明 ROI 之后再分享一部分收益。

主持人：您近一两年与不同国家和企业的一号位交流。最让您欣喜的现象是什么？

李开复：越来越多一号位意识到，他们必须做这件事，而且似乎只有一号位能够作出这个决定。因为这涉及担责，也涉及谁来管理整个企业或国家。我们的转化率并不是特别高，但每一位一号位的认可，往往都会带来很大的项目，也可能真正改变企业的财务表现。

主持人：最大的挫折感又是什么？

李开复：有两种情况。第一种是一号位听完后说“请跟我的 CIO 谈”，团队礼貌对接后，对方只问报价。只关心报价的公司，不是我们最想服务的客户。第二种是 CEO 懂技术，也组织很大的团队与我们反复交流，见了 2 次、3 次、4 次、5 次，后来没有消息了，因为他认为自己已经学会，可以自己做。我判断他可能做不出来。

主持人：企业使用 AI、组织发生变革后，一些员工可能被裁。所谓 AI 素养，是让员工接受被裁的事实吗？员工要避免这种处境，需要哪些条件？

李开复：AI 时代仍有很多机会。首先要学好 AI，不能把 AI 当作敌人，而要把它当作伙伴；其次，最好能成为 DRI。即使在用 AI 解决问题的过程中发现自己原来的价值降低了，也要继续寻找其他价值。解决问题也许比不过 AI，但能不能选择正确的问题？要知道应该解决什么问题，要有商业头脑。你可以做一家一人公司，用 AI 创造收入；如果这些都不合适，还可以转向需要温度的工作和服务业。

今天有些人认为服务业不是顶级工作，但这种认识会改变。历史上，很多职业的社会评价都发生过变化。未来如果 AI 推动经济高速增长，人们不再担心温饱，那些为社会带来温度、获得客户感谢和热爱的人，难道不应该最受社会尊重吗？

主持人：最后一个问题：目前您最希望 AI 帮助解决、自己最感到力不从心的问题是什么？

李开复：一年前，怎样说服 CEO 接受这套理念是比较难的。后来，AI 帮助我们设计了老板 AI 这个产品。今天更大的问题可能是传播：信息不对称仍然很严重。现场很多年轻人可能听懂了，但如果把中国最大的 1000 家企业负责人请来，真正听懂的也许只有 10%。怎样让更多人理解，这是我现在最想解决的问题。

来源：商学院

<https://mp.weixin.qq.com/s/7QN8luIXLQvi6UxcJLTrSw>

（五）探求人工智能的第一性原理：把握硅碳时代的人机共生（香港科技大学首席副校长，中国工程院、英国皇家工程院、欧洲科学院院士郭毅可）

演讲核心观点：

他表示我们正进入一个硅碳时代，硅基生命跟碳基生命将共同塑造我们的未来。通过梳理人类智能和人工智能的科学发现史，郭毅可教授论证了“硅碳智能在物理上是同源的，都是反熵增，建立局部秩序；数学上是同构的，都是在进行贝叶斯大脑的推理”。在这个全新的碳硅共生时代，人类是主导的。共生不是人类退场，而是机器负责人类希望它负责的工作，把人类解放出来，“在目标设定、价值判断、同理心、情感等方面负担使命”。

来源：世界人工智能大会

https://mp.weixin.qq.com/s/P-RIo8_v0TGcYGQIgqX3-w

演讲原文：

大家早上好！非常高兴今天能和大家一起探讨一个极具价值的命题——人工智能的第一性原理，以及我们当下身处的硅碳时代，硅基生命与碳基生命如何共生共荣，共同塑造人类的未来。

在这个技术飞速迭代的时代，我们需要先回答几个最根本的问题：为什么碳基生命和硅基生命能够共存？为什么人类的智能，如今能够在机器上得以延伸和发挥？更进一步，我们常

说碳基生命、硅基生命，那生命的本质究竟是什么？我们谈论碳基智能、硅基智能，那智能的核心又是什么？这些基础问题，是我们开展相关科学研究、把握未来发展方向的前提，只有把这些最根本的问题讲清楚，我们才能对人工智能的研究和未来走向，有更清晰、更深刻的认知。

01

从“人”到“人工智能”的科学奠基人

如果我们要探讨人工智能，就要先理解“人是什么”，有6位科学家的贡献至关重要，他们分别是达尔文、薛定谔、克里克、霍夫、图灵和维纳。

首先我们来谈谈薛定谔和他1943年出版的名著《生命是什么》。提到薛定谔，大家可能首先会想到“薛定谔的猫”，他是量子力学领域的核心人物，但他流传最广、影响最深远的作品，正是这本《生命是什么》。这本书并非单纯的科普读物，其中蕴含着大量原创性的深刻思考，核心就是对生命本质的解读。

薛定谔的核心观点，源于热力学第二定律：任何孤立系统，都会自然地走向熵增，也就是走向混乱和无序。如果将整个宇宙视为一个孤立系统，那么它最终会陷入混沌的热寂状态。但我们放眼大千世界，却能看到无数局部的有序系统，这背后的机制是什么？薛定谔给出了答案：生命的本质，就是对抗熵增。生命并非封闭的动力系统，而是开放系统，它通过与外界环境

交换物质和能量，获取“负熵”，也就是抵抗熵增的能力，从而维持局部的有序状态。这一观点，为后续整个生命科学的研究，提供了核心指引。

既然生命能够对抗熵增、维持局部有序，且生命具有进化和延续的特性，那么这种机制必然需要一种载体来传承，就像计算机程序需要脚本一样。

在薛定谔思想的指引下，10年后，克里克发现了DNA，这就是生命的“遗传脚本”，也是我们的遗传程序。它存在于我们的细胞之中，能够被读取、转录为RNA——我们接种的mRNA疫苗，正是基于这一原理。RNA会执行遗传程序的功能，生成蛋白质，而蛋白质之间的相互作用，就构成了我们细胞的各项活动，这就是分子生物学中著名的“中央法则”。

如果我们将其与机器运行逻辑类比，DNA就像是存储在磁盘上的程序，RNA是读取到内存中执行的程序，而蛋白质就是程序运行后产生的结果。由此可见，生命不仅是碳基分子的组合，更是生命信息的载体，它通过执行信息，实现减少混乱、构建秩序的功能。而在这个过程中，大脑就像是生命系统的CPU和司令部，是进化的产物。我们的大脑中，留存着2亿多年前爬行动物的烙印，中脑源自1亿多年前的哺乳动物脑，这些都是遗传而来的，这也是我们与许多动物拥有相似行为和功能的原因。而最外层的大脑皮层，才是人类高级智能产生的核心脑区，这里布满了神经元。

神经元是一种特殊的细胞，拥有细胞核，外部延伸出轴突和树突，最尖端的部分叫做突触。神经元之间通过突触相互连接，并非直接接触，而是通过离子交换实现信息传递，这就形成了神经网络。我们将不同动物的神经元数量进行对比会发现，果蝇的神经元数量极少，而人类大脑皮层的神经元数量，在所有动物中是最多的，达到 160 亿以上。可能有人会说，大象的神经元数量有 2500 多亿，但大象的神经元主要分布在小脑，大脑皮层的神经元仅有 60 多亿，远少于人类。当然，我们也不能小看大象，它也是一种非常聪明的动物。

了解了神经元的数量，我们再来看看它的工作机制。我们今天所说的人工神经网络，本质上就是对大脑工作模式的模仿。人类的大脑约重 1.3 公斤，拥有 860 亿左右的神经元，神经元之间的连接会传递离子，进而产生脑电活动，这也是我们能够通过脑电仪检测到脑电波、实现脑机接口的原因。这些神经元的连接，是意识的物理实现，而核心问题在于：这些连接是先天固定的，还是可以改变的？也就是我们所说的“可塑性”。

加拿大著名神经科学家霍夫，在 1948 年出版的《行为的组织》一书中，提出了“大脑可塑性”理论，也被称为霍夫学习理论。他指出，当我们的的大脑受到刺激时，部分神经元会同时活跃起来，而共同活跃的神经元，会建立起更紧密的连接，这就是“共同激活的神经元共同连接”。这一理论，奠定了学习的生物学基础：我们通过不断接收外界信息、刺激大脑，能够强

化神经元之间的连接，也能抑制部分连接，从而实现大脑的“重塑”，这就是学习的本质，也是今天所有人工神经网络的核心生物学依据。

前面我们探讨的都是碳基生命的智能机制，而在 1936 年，计算机科学的鼻祖、英国数学家和哲学家图灵，跳出了这一局限，提出了一个颠覆性的问题：思维和计算，必须依赖人类的大脑这一碳基载体吗？答案是否定的。图灵聚焦于思维与计算的本质，发明了“图灵机”这一抽象概念，证明所有的计算，都可以通过一个简单的机器来实现。

图灵机的核心组成非常简单：一是纸带，相当于程序的抽象，存储着指令序列，就像我们的 DNA 是遗传指令的序列一样；二是读取指令的机器，相当于 CPU；三是根据指令改变读取位置、输出结果的模块。这一抽象模型，就是现代计算机的原始雏形。15 年后，美籍匈牙利科学家冯·诺依曼，在图灵机的基础上，制造出了第一台实体计算机，让图灵的理论变成了现实。

与图灵同时代的麻省理工学者维纳，深受薛定谔思想的影响，他提出了一个新的问题：既然生命的本质是抵抗熵增、建立局部秩序，那么这种机制，能否在机器系统中实现？维纳早年是生物学家，后来转向物理和数学领域，学识极为渊博。他认为，生命的核心并非碳基载体本身，而是抵抗熵增、构建秩

序的机制，只要机器能够实现这一机制，就具备了类生命的属性。基于这一思考，他创立了划时代的控制论。

控制论的核心，是通过反馈机制缩小系统的可能性空间，实现熵减。所谓反馈，就是设定一个目标，执行相关行动后，将行动结果与目标进行对比，若存在差距，就修正行动，这也是导弹的核心工作原理。维纳在 1949 年发表了著名文章《人有人的位置》，在这篇文章中，他首次预判了未来人类将与智能机器共存场景，也就是一个能够自动实现熵减的系统，与人类共同构成新的社会形态。

02

信息、先验认知、贝叶斯推理

我们一直在谈论“信息”，那么信息究竟是什么？

我们知道，生命需要吸取负熵来对抗熵增，而负熵的核心，就是能够减少不确定性的信息。在人与人、人与环境、机器与环境的交互中，我们会传递各种“消息”，也就是对某种状态的描述，但消息并不等同于信息。信息的核心价值，在于它能减少接收者的不确定性。

举个例子，假设我们看世界杯，我说“中国队输给了阿根廷”，这个消息对你来说可能毫无兴趣，因为你原本就认为中国队实力不如阿根廷，这个消息没有减少你的不确定性；但如果我说“中国队赢了阿根廷”，这个消息就会极大地减少你的不确定性，对你来说就是有价值的信息。

再比如，我说“人终有一死”，这是一个既定事实，没有减少你的不确定性；但如果我说“我明天就会死”，这个消息就会让你产生极大的震动，因为它大幅改变了你对相关事情的认知。所谓认知，就是通过吸收这些能够减少不确定性的负熵信息，建立起的思维结构。

理解了这些，我们再来看今天的人工智能神经网络，就会发现它的核心逻辑，就是将人类神经元的架构，迁移到计算机和云计算平台上。人类的神经元有细胞核、轴突、树突和突触，人工神经网络也有对应的结构，通过激活、静息状态的切换，实现信息的传递。而人工神经网络中的“权重”，就相当于人类神经元之间连接的紧密程度，权重的调整，就是对大脑可塑性的模仿，也就是“学习”的过程。

根据霍夫学习理论，学习的本质就是改变神经元之间的连接权重。在人工智能中，我们通过给模型输入数据、设定输出目标，让模型不断调整权重，直到输出结果与目标一致，这个过程就是“拟合数据”；调整权重的算法，就是学习算法；而模型的输出结果与目标结果的差值，就是“损失”。

机器学习的过程，就是通过一次次迭代，不断降低损失，让模型逐渐收敛、变得精准。比如，GPT 有 171 亿个参数，本质上就是 171 亿个连接权重，连接数越多、参数规模越大，模型的学习能力和生成表现通常就越好。

讲到这里，我们需要探讨一个更深入的问题：认知的本质是什么？认知是主观的，是每个人对世界的理解，这种理解因人而异，并非绝对统一。这是因为我们对世界有一个“先验认知”。

这就引出了“预测编码理论”，这也是今天人工智能研究的核心理论基础之一。我们对世界的认知，并非被动接收外界信息，而是基于先验认知，对世界进行“猜想”。我们的五官将外界信号传递到大脑，大脑根据先验认知，构造出一个对世界的模型，当外界信号与我们的预测一致时，我们不会产生新的认知；只有当信号与预测存在差距，也就是出现“意外”时，新的信息才会进入大脑，推动我们更新认知。

这种认知模式，是符合生命熵减需求的最优机制。

首先，我们的五官接收的外界信息极为庞大，大脑必须通过先验认知过滤掉噪声，只关注有价值的信息；其次，这种预测模式能弥补信息传递的延迟。比如足球向我们飞来时，视觉信号传递到大脑需要 0.2 秒，若没有预测，我们根本无法接住足球；最后，人类大脑的功耗仅为 24 瓦，这种高效的认知模式，能最大程度降低能量消耗，避免大脑过载。

那么，我们的学习过程，本质上就是一个“信息压缩”的过程。首先，外界的光波、声波等物理信号，经过我们感官的筛选，形成神经信号，这是第一次压缩；其次，神经信号在大脑中转化为神经网络的连接结构，形成对世界的模型，这是第二

次压缩；再次，我们用语言为这些模型命名、描述，语言是人类独有的能力，是对大脑中世界模型的表达，这是第三次压缩；最后，我们将语言描述的规律，凝练为数学公式，这是第四次压缩。人类的智能活动，就是在这样一次次的信息压缩中实现的。

而实现认知的核心过程，就是“贝叶斯推理”。简单来说，贝叶斯推理的逻辑是：我们先有一个先验模型，然后通过观察外界信息，将观察结果与先验模型结合，更新为新的模型。观察结果与先验模型的差距，就是认知的动力。这也印证了一句话：“你不是因为看见了才相信，你是因为相信了才看见。”我们的新认知，并非凭空产生，而是基于先验认知，结合新的观察证据，不断修正、完善的结果。

贝叶斯推理的核心公式非常简单：后验概率正比于似然度乘以先验概率。这个公式，就是今天大数据和人工智能训练的核心底层逻辑。

大数据的本质，就是通过我们能够获取的观察数据，构建对未知世界的模型，然后通过反推、验证，不断修正模型；而人工智能的训练过程，就是不断用新的观察数据，更新模型的先验认知，缩小预测与实际的差距。

我们日常生活中的每一个决策，其实都在遵循贝叶斯推理：面试时，我们先通过简历建立对候选人的先验认知，再通过提

问、交流进行观察，根据观察结果修正对候选人的评价，最终做出决策。

理解了认知的本质，我们再来看大语言模型，就会明白它的核心逻辑。大语言模型的本质，是将语言再次压缩到人工神经元上。它通过“向量编码”，将语义相近的词语转化为相近的向量，用向量来衡量词语之间的相似性，而且这种向量是动态可计算的，比如“苹果”这个词，在水果的语境中是一个向量，在手机的语境中则是另一个向量。

大语言模型通过“预测下一个词”的任务，构建出一个语言的星图，将所有词语作为节点，根据语义关联将它们连接起来。千万别小看“预测下一个词”这个任务，当模型能够精准预测下一个词时，就意味着它已经复刻了语言背后的世界逻辑——因为语言是对世界的表达，能够重新生成语言，就意味着能够还原语言所描述的世界。从这个意义上来说，大语言模型已经突破了图灵测试，实现了硅基设备与人类在语言世界中的共生，我们已经能够与大语言模型进行自然、流畅的交互，甚至无法分辨对方是机器还是人类。

03

硅基智能和碳基智能同源同构、拥有相同的智能特征

硅基设备为什么能够符合生命熵减的本质？这就需要回归到“自由能原理”。

自由能原理（FreeEnergy Principle, FEP）是英国神经科学家卡尔·弗里斯顿（Karl Friston）于 2006 年提出的一种认知神经科学和复杂系统理论。该原理试图用一个统一的数学框架，解释生物系统（特别是大脑）如何通过最小化“变分自由能”来维持自身的生存、结构完整性和智能行为，被认为是“神经科学的大统一理论”之一。

自由能原理给出了贝叶斯大脑的物理基础，它指出，我们的预测与观察之间的差距，等价于物理上的自由能，而生命的所有行为，都是为了缩小这个差距。数学上可以证明，贝叶斯推理是减少资源消耗、缩小预测差距的最优方式。因此，大脑遵循贝叶斯逻辑，并非进化的偶然选择，而是生命对抗宇宙无序化、建立局部秩序的物理特性所决定的。

讲到这里，我们就可以总结出人工智能的第一性原理，也就是智能的基本行为框架。无论是人类的智能，还是人工智能的智能，核心逻辑都是一致的：我们的大脑（或机器的模型）中，存在一个对世界的主观模型，我们通过观察外界、做出预测，当预测与观察存在差距（也就是认知误差）时，会有两种选择：一种是“学习”，也就是正向推理，通过调整主观模型、更新先验认知，缩小差距；另一种是“主动推理”，也就是改变客观世界，让世界符合我们的预测，这就是今天我们所说的“具身智能”。

而要实现主动推理，就需要“世界模型”，也就是能够预测行动后世界状态的模型。比如，我们想要拿起一个杯子，大脑会给运动神经元发出信号，指挥胳膊做出动作，同时通过世界模型预测动作的结果，再通过反馈修正动作，直到成功拿起杯子。今天我们所说的学习、大语言模型、具身智能、世界模型，都可以纳入这个智能框架之中。从这个意义上来说，机器人就是硬件化的智能体，是物理世界中的智能体。

总结来说，智能的基本实现机制，就是向内调整和向外行动：向内调整，就是改变自己的先验认知、优化内部的世界模型；向外行动，就是通过改变客观世界，缩小预测与实际的差距。而到这里，大家会发现，我们已经分不清自己谈论的是碳基智能还是硅基智能，因为它们在本质上是一致的。

04

不要让 AI 替人类寻找生命的意义

碳基生命和硅基生命，都是在反熵增的基础上，建立局部秩序的不同演化路径。如果我们将薛定谔“负熵为生”作为对生命的通用定义，那么碳基生命的演化路径是：从 DNA 发现遗传密码，从神经生物学理解学习机制，进而构建贝叶斯大脑，最终实现人工智能的技术转化；而硅基生命的演化路径是：图灵提出思维可以脱离碳基载体存在，维纳为机器赋予生命的定义，进而构建社交网络、智能系统。二者在物理底层同源，都

是反熵增建立局部秩序；在运行机制上同构，都是遵循贝叶斯推理的信息处理过程。

所以，我们当下正迈入一个人机共生的时代。虽然硅基智能和碳基智能同源同构、拥有相同的智能特征，但人类始终是这个时代的主导者。这种共生，并非人类退出舞台，而是让机器承担我们希望它承担的事务性工作，把人类从繁琐的事务中解放出来，让我们能够活得更像“人”——这正是维纳在《人有人的位置》中表达的核心观点。在人机共生的时代，人类需要聚焦于自己的核心价值：设定目标、做出价值判断、构建同理心、传递情感、承担责任，这些都是人工智能无法替代的人类专属职能。

人类文明的发展，经历过两次重大革命：第一次是语言的诞生，它让我们能够组织思维、传递思想；第二次是文字的诞生，它让我们能够记录知识、传承文明。而今天，人工智能的出现，相当于为我们赋予了一个外在的智能体，能够帮助我们处理信息、完成事务，这将推动人类进入一个全新的、人机共生的伟大社会。

接下来，我也想回应几个大家比较关心的问题。有朋友问：AI有泡沫吗？我的答案是：AI整体作为人类文明演进的重要阶段，不存在泡沫，但细分的技术路线和投资赛道，可能存在阶段性的热度波动。就像语言和文字的诞生，是人类文明的必然，人工智能的发展也是如此。我们在投资时，不能盲目追逐

热点概念，而要理清 AI 的技术发展脉络，明白不同技术在整个框架中的位置，这样才能做出理性的投资决策。

还有朋友问：如果硅基智能体能够参与社会决策、拥有数字财产，我们该如何在法理和制度上界定它的责任边界？这是一个非常重要的问题。我们人类社会的法律、制度、秩序，都是建立在“人类是唯一主体”的假设之上的，而进入碳硅二元社会后，我们需要重构社会秩序，核心就是厘清人机的责任边界——哪些责任可以转移给硅基智能体，哪些责任必须由人类承担。比如在医疗领域，哪怕机器读片的准确率远高于人类，最终的签字权、追责权，在当前的法律和制度框架下，仍然属于人类，因为只有人类能够承担相应的责任。这是我们未来需要长期探索和完善的课题。

还有朋友问我，作为研究人员，我现在有多少时间和碳基生命、硅基生命打交道？我的答案是：我和两者打交道的的时间，都在与日俱增。以前，程序员需要编写代码，而现在，机器已经能够高效地完成代码编写工作，而且又快又好。这时候，程序员的工作就从编写代码，转向了定义任务、验证结果，甚至要求机器给出代码的科学解释。人类的工作，始终在向更高阶的环节迁移，这也印证了人机共生的核心逻辑：机器承接事务性工作，人类聚焦高价值、高创造性的工作。

同时，很多年轻人现在存在焦虑，担心自己毕业之后，原本可以从事的工作，会被 AI 替代。我无法给大家一个确定的

未来，但我可以分享三个应对不确定性的方法，帮助大家做好准备。

第一，一定要学习自己真正热爱的东西，而不是盲目追逐当下热门的就业赛道。因为当下的热门岗位，很可能在未来被AI替代，而你热爱的东西，会成为你最擅长的领域，也会成为你在未来社会中的核心竞争力。

第二，在人机共生的社会中，要不断强化自己的人性素养。要培养批判性思维、跨领域知识整合能力、情感交互能力，构建属于自己的知识体系。这些能力，是人工智能无法替代的，也是人类区别于硅基智能的核心所在。

第三，要对未来保持乐观的态度。不要陷入“被AI替代”的焦虑中，而要看到AI带来的机遇：AI让人类变得更聪明，人类用更聪明的头脑制造更好的AI，更好的AI又反过来推动人类进步，形成正向循环。保持乐观，才能更好地把握未来的机遇。

我始终对AI的发展保持乐观，但也有一些忧虑。我的忧虑并非来自技术本身，而是来自人类自身。当机器替代了我们的事务性工作，我们拥有了更多的空闲时间，但很多人可能会将这些时间浪费在无意义的事情上，比如沉迷游戏，在物质条件满足后，丧失生活的目标。中国有句古话：“生于忧患，死于安乐。”人类是因劳动而进化的，空闲时间的价值，在于让

我们去做更有意义、更有创造性的事情，而不是陷入安逸和懒惰。

最后我想用一句话结束今天的分享：AI 可以为我们做很多事情，帮助我们处理事务、解决“柴米油盐”的烦恼，但请记住，不要让 AI 替我们寻找生命的意义。生命的意义，需要我们每一个人用自己的热爱、思考和行动，去探索、去定义。

谢谢大家！

来源：商学院

https://mp.weixin.qq.com/s/PePIO_OhtVZuoOdcLw2Z2w

三、2026 人形机器人与具身智能创新发展论坛

（一）AI 进入物理世界：具身智能的技术跃迁与产业拐点（AI 战略专家、人工智能经济学奠基人、“云经济学”之父、美国未来产业研究院院长乔·韦曼（Joe Weinman））

演讲核心观点：

他借用信息卓越、智能嵌入产品、改善客户关系、加速创新四种竞争策略框架，指出 AI 让运营卓越、产品领先、客户亲密不再互斥。他以三个“意想不到”的案例展现具身智能的广度：激光除草机、Alcon 眼科手术机、以及全球并行的机器人科研实验室。他还梳理了从物理场址、电力散热、芯片计算架构到脑机接口等对外交互层的完整技术栈，并以“莱特兄弟首飞”作比：今日看似前沿实则只是起点。

来源：甲子光年

<https://mp.weixin.qq.com/s/kv9ZFdxwIT-SjUGzlkivg>

（二）工业多臂具身智能机器人世界模型（飒智智能董事长、CEO 张建政）

演讲核心观点：

他指出，飒智致力于让机器人从“被动执行”跃迁到“主动认知”——理解“为什么做”并预判“接下来怎么做”，并针对当前 VLA 的三大困局，自研了强调“可审计、可预知、可沟通”的 WASM 世界模型。技术上，玄枢通过 Encoder/Decoder 把视觉、点云、力反馈、工艺文件等多模态信息统一映射到物理语义空

间，再解码为可执行、经物理校验的动作序列，并构建云端进化、边端实时技能调度、安全约束“三大闭环”，配合 PDAM 渐进式蒸馏动作模型。依托“数据—模型—本体—技能经验”四位一体飞轮和数十万小时、数百 TB 工业数据、2000 多个场景的打磨，其成果已在头部工厂规模落地。

来源：甲子光年

<https://mp.weixin.qq.com/s/kv9ZFdxwIT-SjUGzlkivg>

（三）共生之地——上海马桥人工智能创新试验区（上海马桥人工智能创新试验区建设发展有限公司副总经理夏吟）

演讲核心观点：

她介绍，上海马桥人工智能创新试验区位于闵行南部，与浦东张江、徐汇西岸、临港新片区并列为上海四大人工智能融合创新载体，也是大零号湾产业主承载区和上海市人工智能与智能制造双特色园区。这里既有 4000 年马桥文化、60 年工业积淀，又有便捷的交通枢纽与轨交规划，总面积 15.7 平方公里，可提供充足的产业用地与自持载体。目前试验区在地企业超 400 家、规上工业产值超 600 亿元，去年获评国家级具身智能零部件中小企业特色产业集群，集群超 150 家企业，覆盖上游核心零部件、中游本体、下游应用全链条。她还强调园区在教育、医疗、商业、人才公寓等方面配套丰富，将持续推动具身智能集群化，从“一棵树到一片林”，打造“智生产、智生活、智生态”的产城共生家园，诚邀企业布局落地。

来源：甲子光年

<https://mp.weixin.qq.com/s/kv9ZFdxwIT-SjUGzlkivg>

(四) 2026 具身智能发展趋势思考 (北航机器人研究所名誉所长、智友·雅瑞科创平台发起人、中关村智友研究院院长王田苗)

演讲核心观点:

他以一张“价值演进图”揭示了具身智能下半场的发展逻辑，行业正沿着“科学家颠覆创新→ToB 或 ToC 产品平台形成→上游核心部件→材料突破→垂类应用→Agent+数字员工”的路径依次轮动，每一次轮动都孕育着机会，也暗藏洗牌。中美在具身智能赛道上呈现鲜明分野，中国在能源、算力、硬件供应链与垂类 ToB 应用上领先；美国则强于大模型底层理论与 ToC 场景的率先探索。

来源：甲子光年

<https://mp.weixin.qq.com/s/kv9ZFdxwIT-SjUGzlkivg>

(五) 具身世界模型新范式，通往具身 ChatGPT 时刻 (银河通用机器人联合创始人、大模型负责人张直政)

演讲核心观点:

他借鉴数字 AI“AlphaGo—GPT-2—ChatGPT—系统化商业闭环”四阶段指出，具身智能不仅要训练强模型，更要解决“好用、易用”才能规模化落地。技术上，银河通用依托物理仿真

基建，用五种互补数据训练大小脑一体化基模，将世界模型与 VLA 结合、解耦环境建模与策略学习，他还分享了近期最新技术成果 WAM-TTT，仅凭无标注人类视频实现后训练部署，与让机器人部署后自适应磨损、终身学习的技术，核心是靠数据闭环飞轮把基模真正转化为能干活的生产力。本次 WAIC 大会现场，银河通用展示了抗干扰做早餐、机器人零售值守、工业物流搬运、迎宾交互四项核心应用场景。

来源：甲子光年

<https://mp.weixin.qq.com/s/kv9ZFdxwIT-SjUGzlkivg>

（六）从炫技到实干·工业具身落地探讨（埃夫特及启智机器人董事长游玮）

演讲核心观点：

直指具身智能“重演示、轻落地”的行业盲区，强调机器人“量销”远比“量产”重要，核心在于攻克多机交互、工艺强耦合等复杂刚需技能；为此，埃夫特与启智 Openmind 提出以 HALO 全模态采集服和动作迁移模型打破数据孤岛，采用“上层数据驱动大模型+底层数理机理模型兜底”的混合架构，兼顾泛化能力与工业级（99.99%）可靠性，并通过覆盖“采-存-治-训-推-用”的全链路工具链大幅降低开发门槛，携手全行业打造标准化通用技术底座，加速推进具身智能落地。

来源：甲子光年

<https://mp.weixin.qq.com/s/kv9ZFdxwIT-SjUGzlkivg>

（七）下一脚，踢出具身时代的安卓（加速进化创始人、CEO 程昊）

演讲核心观点：

他指出，随着本体能力成熟，产业瓶颈已不在本体，而在“能有多少人基于它创造新能力”。机器人不同于可“读万卷书”的大模型，必须在真实世界“行万里路”——机器人踢足球正是最接近真实世界、最小却足够复杂的训练场，胜负即最直接反馈，能持续产生数据、转动“具身飞轮”。他主张具身智能时代需要像 Windows、安卓那样的开放平台，避免重复造轮子；为此在发布旗舰机 Booster T2 的同时推出 Booster Studio 开发平台，让开发者在仿真中训练验证、一键迁移真机，把机器人从硬件产品变为持续创造能力的平台。他判断未来十年本体将成基础配置、开发门槛大幅降低、机器人终将从产品成长为平台，产业最大竞争在于开放生态。

来源：甲子光年

<https://mp.weixin.qq.com/s/kv9ZFdxwIT-SjUGzlkivg>

四、“智序生命·AI 健未来”浦江医学人工智能论坛

(一) 国家卫生健康委统计信息中心党委书记、主任赵韡
致辞

演讲核心观点：

赵韡充分肯定上海医疗 AI 先行先试的标杆成效，详解国家“人工智能+医疗卫生”发展目标与应用方向，以及全国数智健康建设体系布局。他表示，上海医疗 AI 发展领跑全国，在模型评测体系搭建、中试基地建设、临床技术转化等领域硕果频出，同时倡议行业坚守标准筑基、临床普惠、协同融通发展理念，推动医疗 AI 规范创新、惠民利民。

来源：健康上海

<https://mp.weixin.qq.com/s/JXsi5gNFAjoT0-SOclPZLg>

(二) 上海市卫生健康委党组书记、主任闻大翔致辞

演讲核心观点：

闻大翔表示，医学人工智能作为上海人工智能与生物医药两大先导产业的交叉核心赛道，是催生医疗健康领域新质生产力的关键动能。上海将持续夯实制度与基础设施双底座，以民生需求为导向深耕场景落地，聚力打造开放包容、安全可信、协同共生的医疗 AI 产业集群。

来源：健康上海

<https://mp.weixin.qq.com/s/JXsi5gNFAjoT0-SOclPZLg>

五、AI 赋能民生福祉专题论坛

（一）北京电控党委常委、副总经理董永生致辞

演讲核心观点：

作为国企集团，北京电控始终以“科技+产业”双轮驱动服务国家战略，依托数字科技力量深耕民生赛道，为自主可控人工智能产业落地筑牢根基。北电数智正是落实这一战略的核心产业平台，承担着集团在 AI 惠民事业中的责任与使命。

来源：

https://mp.weixin.qq.com/s/w_HRuunzjIvLRAtWn8qZYQ

（二）北电数智首席科学家窦德景

演讲核心观点：

在论坛上阐释了北电数智技术底座的两项核心升级：一是“数算模用”全栈自研体系的深度整合，实现全链路自主可控；二是可信数据服务的持续突破，已有民生应用实践落地。他指出，推进技术底座的自主可控与安全普惠，核心在于让国产算力从“可用”走向“好用”，为 AI 赋能民生筑牢可信根基。技术底座全栈自主可控，是 AI 落地民生场景的底层支撑，更是 AI 惠民真正触达每一扇家门的根基所在。

来源：

https://mp.weixin.qq.com/s/w_HRuunzjIvLRAtWn8qZYQ

六、2026 中国电信人工智能生态论坛

（一）星辰共生 智惠伙伴（中国电信董事长柯瑞文）

演讲核心观点：

柯瑞文指出，中国电信坚决贯彻落实，全面拥抱人工智能，结合企业实际，全面实施云改数转智惠战略，打造“算力、平台、数据、模型、应用”五位一体智能云体系，构建“L”型布局，为客户提供优质、安全、普惠的一体化智能服务。面对人工智能发展的历史性机遇，中国电信坚持应用导向，持续加强“息壤”第一科技攻关，持续深化智能云体系建设，加快构建新型数字信息基础设施，打造体系化安全防护能力，进一步全面深化改革开放，推动服务型、科技型、安全型企业建设再上新台阶。

一是加大关键核心技术攻关力度。中国电信锚定科技领军企业建设目标，坚定不移加强原始创新和关键核心技术攻关。加快推动“息壤”智算平台升级，强化国产算力芯片适配，攻坚异构算力调度、AI 原生 PaaS 等核心技术，突破异构混训、异构混推等关键能力，实现高水平的算力调度和模型训推。前不久，中国电信参与完成的“鹏城云脑”项目获得国家科学技术进步奖一等奖，其中“息壤”平台在算力调度中发挥重要作用。目前“息壤”平台纳管算力已超过 100EFLOPS。持续强化星辰大模型攻坚，坚持“国芯训国模”，加快攻关大模型核心技术，创新模型架构设计，提升多模态理解、端到端代码生成、智能体自

主决策等关键能力。中国电信建设“国芯训国模”智算集群，大集群下 AI Infra 能力显著提升，实现国产算力环境下大模型训练与演进关键技术的突破。推动 AI 与通信技术深度融合，创新构建智传网体系（AI Flow），攻关家族式同源模型、生成式传输等关键技术，破解远距离、跨空域极限场景的算力传输难题，大幅提升智能服务覆盖范围。智传网荣获了本次大会的最高奖项“卓越人工智能引领者奖”。

二是加快 AI 规模应用。中国电信着力推动 Token 经营，今年 3 月，中国电信率先提出 Token 经营，经过一段时间实践，中国电信进一步加深认识，五位一体智能云体系就是 Token 经营体系，Token 经营本质就是为客户提供 AI 服务。以“智云上海”的实践为例，通过创新打造一站式 AI 服务门户 AI STORE，为广大客户提供便捷优质的算力、模型与 AI 应用，5 月发布以来用户迅猛增长，有效带动 Token 消费快速增长。中国电信着力打造 AI 原生产品，今天中国电信将发布的星辰超级智能体 TeleAgent，在内部试用后得到了基层员工的广泛好评，目前注册超 30 万，日均 Token 消耗超 2000 亿。中国电信持续推进传统产品 AI 升级，打造天翼智铃、智屏、智看、云智手机等核心 AI 产品，不断创新软硬一体的新型 AI 应用，为千家万户带来智能服务新体验。中国电信加快推动 DICT 向 AICT 转型，深入推进“AI+”行动，聚焦政务、交通、制造等行业，深入客户核心生产流程，通过智惠工程师队伍，提供驻场服务，

解决 AI 应用落地最后一公里问题，助力提质增效。

三是构建全方位安全防护体系。“十五五”时期，中国电信将全面发力安全型企业建设。打造智能云体系全栈安全防护能力。围绕算力、网络、数据、模型、应用等安全防护对象，构建覆盖全域的智能协同防护体系和一体化安全运营体系，建优阡陌高质量数据集，打造业内一流的安全·见微大模型，实现威胁自主发现、闭环处置的 AI 内生安全，以“AI 对抗 AI”。打造安全可控的 AI 产品服务。坚持“AI for 安全”“安全 for AI”，推动产品服务向主动防御、智能协同方向演进，提供全栈化安全解决方案。中国电信把安全作为产品服务的内生要求，坚持 AI 服务延伸到哪里，安全保护就同步覆盖到哪里。比如中国电信打造代码安全审计智能体，深度嵌入代码开发工具，上线前即可自动审计，提前发现代码安全风险。中国电信积极参与人工智能治理框架与规则制定，协同产业各方共同研判、积极应对人工智能应用风险，推动人工智能朝着有益、安全、公平的方向健康有序发展。

四是加快建设新型数字信息基础设施。中国电信将适度超前建设算力基础设施。按照“高功率密度、高 IT 产出率、灵活建设、灵活扩展”的标准建设新一代 AIDC。围绕八大枢纽节点持续优化 AIDC 布局、加大储备，目前已建数据中心超 900 个，总电力容量近 8G 瓦，是全国最大的 AIDC 服务商。加大国产算力规划与布局，围绕重点城市强化边缘算力池部署，提升单

节点算力规模与服务能力。中国电信将加快打造弹性敏捷的算力互联网络。建设“新八纵八横”骨干光缆网，加快推进传输链路向 800G 演进；深化千兆光网普及与 5G-A 部署，提升“入算”接入能力。探索超大容量、OCS 全光组网算内网络建设，支持超大规模集群，提升训推算效。中国电信将持续加大海缆资源投入，为人工智能的国际合作提供“海陆互备、内外协同、安全可靠”的网络基础。近日，中国电信自主设计的首艘国产深海型智能化海缆施工船——“天翼领航者”号，在上海正式入列并圆满完成全系统海缆施工试验，实现海底光缆的高效建设与维护，确保算力中心之间数据的高效流通和 Token 跨域流动。

五是进一步全面深化改革开放。中国电信将改革开放作为关键一招，继“光改”“云改”之后，以“智改”为核心，进一步全面深化改革开放，提升企业治理体系和治理能力，加快推动人工智能发展。中国电信将坚持以客户为中心，升级云中台为 AI 中台，推动企业主流程变革，深化集成体系和云网体系改革，强化智惠工程师队伍建设，系统提升面向客户的需求洞察、数据服务、敏捷交付以及持续运营能力。中国电信将进一步加大开放合作力度，围绕智能云体系搭建全面合作生态，持续深化科技、资本与业务等重点领域开放合作，集聚优质资源，携手合作伙伴共同推动人工智能创新发展。落实中央企业“场景万象”专项行动要求，进一步扩大场景开放，共创场景发展。坚持开源开放，推进开源社区建设。依托 WBBA 等国际交流

平台深化技术交流和产业合作，助力弥合数智鸿沟。

来源：人民邮电报

https://mp.weixin.qq.com/s/mAOkY_994Mp_zrAqm0hL0A

（二）专题分享（中国电信首席技术官、首席科学家李学龙）

演讲核心观点：

李学龙现场分享了智传网（AI Flow）面向“天、空、地、海”全域智联互通的前沿创新成果。基于生成式智能传输创新范式，TeleAI 将智传网（AI Flow）解码模型推上卫星；在国内首次攻克了轻小型无人机直连卫星进行高清视频实时稳定传输的难题；把端到端时延压缩到 20-50 毫秒，首次让机器人远程控制范围覆盖全国；首次实现在自然水域远程无线视频传输。现已在中国电信视联网超过 1 亿路摄像头推广落地；服务触达 400 万终端用户；接入千余套应急设备，投入全国一线抗洪救灾。智传网（AI Flow）让“用发信息的带宽打电话、打电话的带宽传视频”成为可能。

来源：中国电信

<https://www.chinatelecom.com.cn/ct/news/lddt/167798.html>

七、AI 影视专场论坛

(一) 通用世界模型——连接数字世界与物理世界的新基础设施 (生数科技创始人、清华大学人工智能研究院副院长朱军)

演讲核心观点:

朱军表示,人工智能正在从理解和生成数字内容,进一步走向对世界规律的建模、推理、预测与行动。通用世界模型将成为连接数字世界与物理世界的重要基础设施,推动 AI 从“生成世界”迈向“理解世界、预测世界、行动于世界”。

来源: IPO 早知道

<https://mp.weixin.qq.com/s/iZXf77UchbortDmDs2AyRw>

演讲全文:

尊敬的各位嘉宾,大家上午好!

首先,非常感谢大家百忙之中参加我们的论坛。我们非常荣幸能够与万兴科技一起,在世界人工智能大会举办这场面向未来的论坛。

大家都知道,大模型正从语言模型走向多模态模型,再进一步走向世界模型,技术能力不断突破。同时, AI 也在从单一工具逐步发展为支撑产业生产力的基础设施,这正是我们发起今天这场论坛的重要初衷。接下来,我会用一点时间,跟大家简要分享一下我们所看到的基础模型发展趋势。

首先,从原理上来看,什么是世界模型?

我们先来看人的世界模型。这里有一个简单的定义。大家都有过骑自行车或者开车的经历，在做出动作之前，我们会在大脑中形成一个小模型，帮助我们理解环境、想象不同动作可能产生的结果、规划行动，最终指导动作输出。这就是人的世界模型，它能够帮助我们适应复杂环境，并作出安全、可靠的行动。

今天，全世界都在关注如何构建机器的世界模型。机器世界模型相关的研究很多，前沿进展也非常快。但从本质上来看，机器的世界模型主要需要具备三项能力。

第一，对环境进行精准理解；第二，能够基于当前状态和即将采取的动作进行预测和想象；第三，能够在预测过程中学习并采取行动。总体而言，我们认为机器世界模型至少要实现理解、想象和行动的一体化。

从这里可以看出，世界模型与普通模型的区别在于，它不仅要回答“是什么”，实现对世界的理解，还要回答“为什么”，以及“这样做会产生什么后果”。这个领域目前最大的变革，是从专用模型走向通用模型，并形成 **Scaling Law**。包括视频模型和世界模型在内，大家都希望遵循 **Scaling Law**，推动通用能力取得重大进展。

要做通用模型，首先需要回答一个问题：我们应该使用什么样的数据？根据理解和建模世界的目标，我们梳理了相关数据，形成了六层数据金字塔。从底层向上，分别包括通用互联

网视频、第一视角的人类动作视频，以及机器人相关的操作数据等。

在机器人轨迹采集过程中，大量数据也是以视频形式承载的。机器人通常配备多个相机，用于拍摄环境和机器人的动作。总体来看，视频是这套数据体系中最主要的格式，因为它是最便捷、最容易记录世界变化的数据载体。

数据金字塔的底层具有海量的数据规模，而在机器人专用操作数据方面，如何收集足够多的数据，一直是行业面临的巨大难题。

我们的目标是打造通用基模，因此在数据使用上有几个原则：第一，数据规模要足够大；第二，数据类型要尽可能全面；第三，要尽可能使用真实数据。这也是我们从大模型训练中总结出的规律。过去，这些数据为什么没有被充分利用起来？一个重要原因是，普通视频与机器人动作轨迹之间缺少直接对应关系。

今天，我们想和大家分享的是，如何通过生成式方法和底层技术原理，让海量视频数据转化为高质量、有价值的训练数据。因此，我们看到，世界模型首先会在数字世界中成长，但最终目的不只是生成数字内容，还要进一步走向物理世界，在更加广泛的场景中发挥价值。

有了前面的数据，接下来要考虑的是，如何真正发挥这些数据的价值。

我们的目标是打造通用模型，而通用模型需要通用架构。简单地将不同模块串联或并联起来，并不是最理想的方式。我们希望在统一架构中，实现理解、想象和行动的一体化。我们较早开始探索这一方向，并提出了全球首个 MoT 统一架构，将理解专家、生成专家和行动专家统一在同一个模型中，通过统一目标进行训练。

因为它是一个一体化模型，所以在使用过程中可以自由切换不同的输出形态。对于机器人，它可以直接输出动作；对于创作者，它可以直接输出高质量的像素级视频；也可以根据任务需要进行联合输出，在一个模型中统一完成不同任务。生数科技正在做的就是这样一件事情「通用世界模型」。

上面展示的是两类模型：世界生成模型与世界动作模型。一类与今天的论坛直接相关，面向数字内容生产和高质量影视生产；另一类则是为具身智能行业提供机器人的“大脑”。

在数字内容方向，对应的是 Vidu Q3 等模型。目前，我们已经取得了很多成果，也推动了内容生产力的提升。

对于一家做基模的公司来说，我们最大的体会是，通过海量视频数据的大规模训练和 Scaling Law，模型可以实现更加精准的理解，以及高动态、高一致性的生成创新，而且这些创新往往是组合式的。这些能力在过去的传统方案中很难实现。如今，在视频大模型中，我们已经率先实现了理解和想象，也就是前面所说世界模型三项能力中的前两项。

近期，我们还发布了实时生成模型。视频模型不再只是简单地生成素材，也可以成为实时交互模型。以数字人场景为例，传统方式通常需要预设有限的动作模板，驱动面部表情和口型。现在，我们采用的是完全流式的生成方式，角色的表现可以更加自然、灵活，也更加拟人化。

它不仅可以在实时交互，还能够支持无限时长的连续对话。这也打开了很多想象空间。在需要流式、实时交互的场景中，可以直接使用视频大模型提供相关解决方案。

当前，流式视频生成主要面临几方面的局限。

第一，缺少实时的用户干预；第二，缺少语音显示控制。语言对人来说是一种非常自然的交互形式，特别是在实时交互场景中，语言是非常重要的输入形态；第三，长时间生成的稳定性通常不足；第四，推理和部署能力不足。对于流式生成而言，无论是推理时延还是推理成本，都有非常高的要求。

针对这些问题，我们进行了一系列技术突破。模型可以端到端地理解语音指令，实现实时交互和指令跟随，完全通过对话方式，还原人类最自然的交流形态。用户也可以定义自己的角色，通过设定人物形象、音色和相关描述，打造具有特定人设的实时交互角色。目前，模型可以实现 540P 的视频输出，最高帧率达到 42FPS。

但我们的最终目标，并不只是停留在数字内容生成上。无论是前面提到的 Vidu 视频生成，还是 Vidu S1 的流式生成，

都是在为进一步走向实时的物理交互做准备。这些核心技术也可以支撑模型在物理世界中的行动。

这也是我们今年 4 月份发布的最新模型。它可以基于同一个模型底座，同时驱动多种异构机器人本体，完成复杂的长程任务。

我们取得的一个重要进展是，基于较强的通用基模，可以用更加高效的方式，让不同机器人本体快速学习复杂技能。这里展示了一些例子。用户可能只需要给出一条文字指令，比如让机器人完成插花、浇水等任务。模型可以理解语言描述，将任务分解为一系列步骤，并精准生成完成任务所需的动作指令，最终控制机器人完成整个任务。

画面中展示的是机器人的“大脑”，也就是我们的基座模型正在预测未来状态。可以看到，模型预测的状态与真实发生的状态非常接近。这说明机器人不只是能够执行任务，还能够对未来进行预测和想象，并基于这种想象，更好地规划和输出动作。

这也是新一代基座模型与传统 VLA 方法最本质的区别之一。传统 VLA 更多是在拟合轨迹：给模型演示一定数量的轨迹，让机器人进行模仿，并在一定范围内实现泛化。通过生成式预训练机制，机器人可以基于对未来的想象来规划动作，输出的动作也更加可靠、稳定、安全和高效。同时，通用基模可以支持跨场景、跨任务的泛化。

图中展示的是传统 VLA 范式的对比结果。蓝色曲线代表我们去年 12 月开源的版本，绿色曲线代表最新采用更强视频模型的版本。图中的结果越靠近左上角，意味着模型可以使用更少的数据，学习到质量更高的任务能力。可以看到，最新版本的提升非常显著。

这意味着，通过通用基座模型，可以大幅提升机器人学习任务的效率和成功率。相关模型在多个榜单中也处于领先地位，特别是在跨场景、跨任务的泛化方面，能够达到非常高的成功率，取得了行业领先的模型效果。

我们今天讨论的通用世界模型，不再只是视频模型的升级，也不是语言模型的简单延伸，而是推动 AI 从生成内容，迈向理解世界、推演世界，并与世界实时交互的新范式。

最后，我想简单介绍一下生数科技在这一过程中的发展历程。

生数科技从成立之初就定位于基础模型，并以视频建模作为出发点。在这个过程中，我们深刻认识到，视频模型基于对世界的深层理解，可以达到非常高的能力上限，而且这个上限仍在不断拓展。这也是我们一直坚持探索的方向。我们希望通过高质量的基模，更加高效地服务高质量内容生产和具身机器人的产业落地。

过去短短两年时间，大家可以看到，视频基模已经从最初展示技术可能性，发展到今天真正进入产业、支撑生产力场景。这是一个非常快速的变化过程。

最后，对于未来的展望...

今天和大家分享的是基础模型的故事。生数科技定位于基础大模型，希望以更高的效率和更高的质量，为合作伙伴和整个行业提供模型能力。

我们的第一个体会是，视频数据是语言之外一种非常丰富、规模庞大且十分宝贵的信息载体，能够记录和描述这个世界。

第二，基于视频的大规模预训练，模型可以实现精准理解、高质量和高一致性的想象与预测，并进一步形成可泛化、跨任务、跨本体的行动能力。这也是我们一直致力于探索的方向。

我们已经看到，在持续提升预训练能力的基础上，再配合高效的后训练，可以推动整个行业实现快速迭代和落地。

这也是我们一直致力于打造的技术范式。希望能够与在座的各位，以及更多行业同仁开展合作，共同探索通用世界模型的智能上限与产业落地的更多可能。

我的分享就到这里，谢谢大家！

来源：IPO 早知道

<https://mp.weixin.qq.com/s/iZXf77UchbortDmDs2AyRw>

（二）万兴科技创始人、董事长吴太兵

演讲核心观点：

吴太兵提出 AI 影视“生产范式重构”的核心产业判断，他强调 AI 影视的爆发背后是创作者经济的重构，AI 绝非影视行业的效率辅助工具，而是正在彻底重构产业运行的底层逻辑，影视产业正站在第四次技术跃迁的关键起点，AI 影视是具备千亿级增长潜力的全新数字内容赛道。他指出，全球影视竞争的下一阶段，不只是内容竞争，更是全球产业生态的竞争。同时现场发出倡议，当前通用世界模型技术持续迭代，AI 影视产业迎来爆发窗口，中国数字内容出海进入全新黄金周期，全产业链伙伴可携手同行，共赴 AI 影视出海大时代。

来源：短剧新圈

<https://mp.weixin.qq.com/s/Ty40RoVkX8MPVk3Bx3fDFQ>

（三）阿里云智能集团资深副总裁、公共云事业部总裁刘伟光

演讲核心观点：

刘伟光表示，世界模型开启新一轮生产力革命，对算力底座、训推一体、数据与工程化能力提出挑战；阿里云提供从数据、算力、训练、推理到 Agent 部署及全球化的全栈支撑，为世界模型迭代全栈就绪。

来源：短剧新圈

<https://mp.weixin.qq.com/s/Ty40RoVkX8MPVk3Bx3fDFQ>

华信咨询设计研究院整理

八、启明创投·创业与投资论坛

（一）通用世界模型：打通虚实边界，筑基物理智能（生数科技联合创始人、首席执行官骆怡航）

演讲核心观点：

AI 正在从理解与生成数字世界，走向理解并行动于真实的物理世界，世界模型正是这场变革的核心基座。生数科技是全球首个实现数字世界和物理世界统一的通用世界模型公司，构建起从理解到生成到行动的通用世界模型，为数字生成和物理行动提供智能。

来源：启明创投

<https://mp.weixin.qq.com/s/BkTiVf0IehylOtA0QrXx7A>

（二）AI 基础设施与模型能力的融合催生（阶跃星辰联合创始人、首席技术官朱亦博）

演讲核心观点：

随着 Token 需求快速增长、模型规模持续扩大，仅依赖算力扩张难以为继，行业需要推动模型、系统与芯片协同设计，在单位算力上释放更多智能。面向智能终端，Infra 还需进一步支持端云协同、长期运行环境等，使端侧实时执行与云端复杂推理高效配合，成为 Agent 真正规模化落地的关键底座。

来源：启明创投

<https://mp.weixin.qq.com/s/BkTiVf0IehylOtA0QrXx7A>

九、“智见 AI 合创未来”人工智能投融资论坛

（一）中国人工智能芯片需要原创技术（国际欧亚科学院院士、清华大学教授、东方算芯董事长魏少军）

演讲核心观点：

魏少军表示，若以人的身体作为类比，农业革命解决了躯干的生命维持问题，工业革命延伸了四肢的能力，信息革命扩展了五官的感知，而正在发生的智能革命则是在延伸人类的大脑，即认知能力。而算力作为人工智能的技术底座，没有芯片就无法实现认知能力的延伸。

来源：明亮公司

<https://mp.weixin.qq.com/s/V7v6Q7bwJ4fPxYdYxjrWnA>

（二）海通国际证券集团有限公司首席经济学家张忆东

演讲核心观点：

A 股具备政策、盈利与资金共同驱动的价值重估基础，左侧布局股市秋季行情的时机呼之欲出。本轮 AI 科技浪潮尚未结束，但已经由“修路铺网”的上半场逐步进入应用扩散和商业化加速的下半场。AI 板块经历 6 月、7 月调整之后，有望启动秋季行情。

来源：中国证券报

<https://mp.weixin.qq.com/s/1mDklbTcxMxjflK2J0aUog>

（三）国泰海通业务总监郁伟君

演讲核心观点：

郁伟君在演讲中表示，科技创新成为境内外资本市场主线，AI 产业企业资本运作正当时，建议优质企业在成长期启动 IPO 准备工作，积极申报 A 股双创板块。国泰海通为 AI 产业企业提供全生命周期综合金融服务，期待助推更多 AI 产业企业登陆资本市场。

来源：上海金融官微

<https://mp.weixin.qq.com/s/WAYkGV-QwXi4mHV98k3ML>

A

（四）全链协同 共创 AI 黄金时代（上海联和投资有限公司董事长、上海兆芯集成电路股份有限公司董事长叶峻）

演讲核心观点：

01 智能体时代全面开启

算力范式迎来关键切换

当前，人工智能正经历第四次发展浪潮。随着生成式 AI 和智能体的快速发展，AI 开始具备“感知-决策-执行”完整闭环能力，能够自主拆解任务、调用工具、迭代优化，进化为能够独立执行任务流程的“数字员工”。这一质变推动 AI 全面迈入行业规模化落地新阶段。

“从训练为王到推理和智能体的范式切换，让 CPU 的战略价值上升到了前所未有的新高度。”叶峻强调。在智能体架构中，CPU 承担着复杂逻辑调度、任务流编排、系统交互等核心职能，CPU 与 GPU 的配比正从训练时代的 1:8 向智能体时

代的 1:1 加速逆转。与此同时，AI 应用正从云端快速向端侧渗透，AIPC、汽车、工业终端等海量场景对算力、经济推理、数据安全提出了刚性需求，CPU+GPU+NPU 异构计算，将成为兼顾成本、效率、规模与场景适应性的更优解。

02 ZX86 自主算力筑基 面向行业场景 共建全栈 AI 方案

兆芯是联和投资重点布局的集成电路设计公司，经过十余年的发展，已成长为国产高性能 CPU 领域的领军企业。兆芯完整掌握指令集、内核微架构、高速互连等 11 项关键核心技术，建立有完善的 ZX86 指令集方案，通过系统构建 ZX86 自主指令集知识产权体系，走出了一条"自主不封闭、兼容不依附、面向新应用"的特色发展道路，可确保后续 CPU 产品及应用全面自主迭代创新，并面向 AI 驱动的行业新场景、新应用、新安全等发展需求提供安全、高效的核心算力。

面向 AI 时代持续增长的算力需求，兆芯以开胜 KH-50000 服务器处理器、开先 KX-7000 端侧处理器等代表产品为基础支撑，与来自整机、国产 GPU、AI 加速器、大模型等产业链伙伴深度联动，协同构建面向全行业的 ZX86 全栈 AI 自主计算解决方案，从底层芯片到上层智能政务问答、工业质检、医学影像辅助诊断、金融实时风控等真实场景应用，整条产业链互相赋能，共同推进千行百业数智化转型升级。

03 全链协同共振

共绘中国 AI 产业新蓝图

当前中国 AI 产业正处于历史性的战略机遇期。“十五五”规划明确将“全国一体化算力网”纳入国家六大新型基础设施，提出 AI 基础设施建设“规模化、集约化、绿色化、普惠化、自主可控”五大原则，为产业发展指明了方向。AI 新基建涵盖硬件、软件、服务全产业链，拥有万亿级的广阔市场空间，在技术迭代与需求升级的双重驱动下，AI 产业具备极强的长期增长确定性。

作为国产自主算力的核心企业，兆芯致力积极发挥产业领军作用，大力加强异构集成、AI 协同、硅光互连等关键技术创新研发，加速推进 ZX86 CPU 产品升级，联合上下游合作伙伴，共同打造从芯片到整机、从硬件到软件、从技术到场景的全栈 ZX86 自主生态，以算力自主支撑产业自主，以技术创新驱动价值创造。

来源：兆芯

<https://mp.weixin.qq.com/s/2urdyA6vV31yHyBs85EnIA>

十、AI 算力技术架构创新论坛

(一) 中国电子信息产业发展研究院党委副书记、人工智能工委会长胡国栋致辞

演讲核心观点：

随着 AI 基础设施加速演进，多元算力之间的互联互通与生态协同，正成为影响算力效能释放的重要因素。面向这一产业趋势，标准建设成为推动技术创新、产业协作和规模化落地的重要支撑。人工智能是新一轮科技革命和产业变革的重要驱动力量。面向人工智能技术快速演进与产业规模化发展的关键阶段，业界需要进一步凝聚共识，坚持需求牵引、应用导向和前瞻布局，扎实推进人工智能芯片重点标准研究与应用验证，更好支撑我国人工智能芯片产业和人工智能产业高质量发展。

来源：中国电子信息产业发展研究院

https://mp.weixin.qq.com/s/CH82AxY_N6WzfweL1erFGQ

(二) 新华三集团高级副总裁、技术委员会副主席刘新民

演讲核心观点：

刘新民指出：当前 AI 竞争的焦点已从单纯追求硬件规模，转向系统级效能与生态协同的综合竞争。新华三秉持多元开放生态理念，围绕异构算力、算力互联、融合软件、整机系统和产业模式等多个维度协同推进：通过兼容多平台主流 xPU 芯片，打造灵活可选的异构算力底座；通过开放互联协议与算网深度融合，提升多元算力的连接效率和协同能力；通过算力资源融

合与统一调度，打破资源孤岛，提升应用适配与部署效率；通过整机系统标准化和模块化设计，推动 AI 基础设施实现更高效率、更高可靠和可持续演进。未来，新华三将携手产业链伙伴，共同构建开放协同的下一代 AI 基础设施生态。

来源：中国电子信息产业发展研究院

https://mp.weixin.qq.com/s/CH82AxY_N6WzfweL1erFGQ

十一、端侧 AI 创新和行业发展论坛

(一) 智能体 AI 时代：计算架构与模型效率的持续创新
(高通公司全球副总裁、中国区研发负责人、IEEE Fellow 徐皓)

演讲核心观点：

徐皓博士指出，端侧的 AI 应用拥有非常广阔的前景，并系统性阐述了智能体浪潮对终端计算架构带来的范式变革。他指出，随着手机使用场景从单轮对话迈向多轮交互、再迈向智能体调度，设备所需的 Token 处理量正以十倍量级递增——从一万、十万到百万级，因此端云协同处理十分重要。为此，高通针对持续运行始终在线工作负载、持续的传感器数据处理和系统性协同需求，面向智能体 AI 打造新一代设备架构，并与面壁智能等合作伙伴持续推动在端侧模型领域的合作。

来源：雷锋网

<https://www.leiphone.com/category/chips/dX9rUFVHGEWJz4yQ.html>

演讲全文：

非常高兴和大家分享高通在端侧 AI 方面所做的硬件、软件、算法和系统开发以及生态合作。

这支视频带我们简单回顾了我们使用手机这一终端的方式，从最开始简单的交互，到今天支持智能体的手机。同样，能够运行智能体的终端品类多种多样，从最小的智能耳机到千

瓦级的数据中心，高通不同的硬件平台可以为这一系列的智能终端提供支持。很早之前，我们就注意到面壁在端侧的布局，我们双方有很多的契合点。

接下来，我们看一下为什么端侧这么重要。这里呈现的是一个简单的统计，可以看到不同终端的市场规模在几十亿量级。大多数人现在都很注重数据中心这类大规模训练，但实际上端侧的 AI 应用也拥有非常广阔的前景。

可以看到，手机从简单的单轮对话、到多轮对话，再到运用智能体，其所需的 Token 量是以十倍量级往上增加的，从一万、十万到百万级。比如，在手机端使用龙虾或者智能体要用到的 Token 数量非常之大，现阶段是端云结合的解决方案，端侧小模型处理隐私性强需要实时处理的任务，云端处理深度推理和规划。将来越来越多的模型处理会在端侧来减少对数据传输和云端 Token 运算的需求。

刚刚面壁的同事也介绍了，端侧模型在未来一至两年内，可以达到与几十倍算力的云端模型同样的功能，并且这个趋势还在不断发展，越来越多的大模型将迁移到端侧，以小模型的形式出现。随着硬件算力的不断提高，这些模型在端侧也将能够轻松运行。

我们看一下，从端侧硬件设计的角度来看，有哪些最重要的设计难点：

第一，在打造智能体手机时，我们需要出色的任务规划，我们通过全新的 CPU 来做任务规划和控制。

第二，需要随时随地处理传感器信息。以前的手机只是用户给它指令让它执行任务，它不需要 24 小时判断你在干什么或者你需要什么。但是智能体手机需要主动帮助你、告诉你需要做什么，需要随时感知你在什么地方，比如它感知你现在在大会现场，它可以帮你自动静音；比如它感知你进入座舱，它会自动调节成你喜欢的温度和座椅位置。这都需要有出色的低功耗传感器处理。传感器中枢会对大量的传感器做有效地处理和融合，来精确地感知环境，给智能体提供决策信息。

第三是 NPU，来提供强大的模型运算功能。我们会做一系列的模型适配和优化，有效地利用 NPU 高效张量计算能力来处理 2B 到 3B 量级的模型在手机端侧的运行。

此外，还需要有出色的 5G 或 6G 连接。现在大部分智能体的功能是在云端运行的，因为需要大量算力，但随着越来越多的模型在端侧运行，需要端侧、云端和边缘云的协同处理，出色的无线连接功能可以把端侧的算力和云端的算力无缝连接。

接下来让我们看看，高通如何带来最有效的硬件支持。首先是高通传感器中枢，它能够为智能体提供更多数据，24 小时在线了解关于你的所有信息，为大语言模型或端侧模型的输入提供情境感知，打造个人知识图谱。另一方面，我们的 NPU

可为端侧提供强劲算力，例如刚才提到的 20 亿、30 亿参数 Min iCPM 大模型，都能在端侧流畅运行。

这是高通最新的旗舰移动平台第五代骁龙 8 至尊版，今年我们还将推出更强大的移动平台，来支持所有的端侧运算；我们的座舱芯片也具备强大算力，可支持更大规模的模型，在车上稳定运行百亿以上的参数大模型。

从应用以及模型的角度来说，我们已经完成从最初的感知 AI，到现在的生成式 AI 和智能体 AI 的迭代，下一阶段我们将重点推进物理 AI，实现从小模型向多功能大模型的转变。我们在算法方面也做了一系列的模型在端侧的适配，比如模型压缩和量化，长上下文的处理和增强，以及端侧模型的不断学习（on device learning）。

最后是从车到机器人的演进过渡，一方面我们要将二维的规划，升级为对三维真实世界的理解；另一方面是针对更多、更复杂的应用场景进行优化与设计。这是我们从汽车芯片到机器人芯片领域需要开展的一系列开发工作。

高通与面壁在手机和车等方面已经有很多深度合作。我们也非常期待未来联合面壁及在座的各位合作伙伴，共同推进智能手机和多种智能终端在端侧的应用和发展。谢谢大家。

来源：雷锋网

<https://www.leiphone.com/category/chips/dX9rUFVHGEWJz4yQ.html>

（二）AI 手机时代来临，未来 3 至 5 年内或将迎来端侧 AGI （面壁智能 CEO 李大海）

演讲原文：

2026 年，随着端侧大模型落地，属于 AI 手机的时代来了。

就在本届 WAIC 开幕两天前，国家网信办公布了 7 款提供手机端侧生成式人工智能服务的备案信息，Apple 智能、华为小艺 AI 大模型、小米澎湃 AI、三星盖乐世 AI、OPPO AndesGPT、vivo 蓝心端侧大模型及努比亚豆包手机大模型悉数在列。

两大国际品牌中，Apple 采用了千问和百度作为供应商，而三星选择了面壁智能 Min iCPM 系列作为端侧模型能力提供方，深度参与此次备案落地。与此同时，据多方信息披露，面壁智能已完成新一轮融资，本轮融资引入中网投、国新基金、建投投资等国家级产业基金。公司 2026 年上半年累计融资金额超 50 亿元，投后估值超 200 亿元，一跃成为端侧智能领域公开估值最大的独角兽企业。

当下，囿于成本、时延、隐私三大约束，大模型正经历第二次结构性跃迁：从以云端为唯一算力中枢的单点智能，走向端云分工、协同共生的分布式智能系统，端侧 AI 已成为产业焦点。

面壁智能 CEO 李大海在论坛致辞中表示："2026 年是端侧 AI 规模化落地的元年。面壁智能的定位很清晰——做普惠

千行百业的大模型基座：以技术引领做深底座，以产业推动做实场景，以生态共建做大未来。”

谈及行业发展以及未来可能面临的竞争，李大海表示，如果一个产业链只有一家公司在做，那可能是在“逆行”。“越来越多合作伙伴和友商加入这个赛道，反而论证了我们前瞻性的眼光与过去看法的正确性，这是一件好事。”

李大海以移动互联网时代的发展作类比，认为当前 AI 的发展阶段，与 2012 年至 2013 年的早期移动互联网阶段高度相似——技术正处于快速发展期，亟待与千行百业进行深度融合，整体市场空间极为广阔。“一家公司在如此庞大的领域中进行创新，其带宽是有限的；而众多伙伴的加入，能把产业的‘创新带宽’拓得非常宽。从产业角度来看，这能让我们更快地找到真正有价值的产业、场景以及有效的落地方式，从而让整个产业更快受益。”

在面壁智能的研判中，当前的手机端侧智能化处于早期阶段。可作为参考的案例是智能汽车。“在如今相对稳定的状态下，自动驾驶技术呈现出 70%由第三方提供、30%自主研发的格局。我们推测，随着手机产业的发展，未来可能也会出现类似的情况。”这背后的底层经济学规律是“社会化分工”——让擅长的公司把事情做好并提供给客户，从而实现边际成本最低、全局产出最优。

至于人们何时能迎来端侧 AGI，李大海预计在未来 3 至 5 年内，时间晚于云端。由于端侧落地需要在有限的算力、有限的带宽以及受限的散热功耗下推进，在具体推进时，需要与厂商和芯片进行更多的磨合，因此天然地需要花费更长时间。相较于云端，端侧爆发的时间会相对延后，这是客观规律。

来源：明真 AI

<https://mp.weixin.qq.com/s/v6663md7Lgax57n5cJrHsg>

十二、人工智能终端产业演进与协同发展论坛

（一）中国联通副总经理郝立谦出席论坛并致辞

演讲核心观点：

郝立谦指出，当下消费电子终端正迎来颠覆性变革，AI从附加功能变为终端核心内核。中国联通发挥 5G-A 泛在连接、全域分布智能算力、脱敏网络数据资产、亿万级线下渠道四大核心优势，既是打通 AI 终端规模化落地的基建支柱，也是串联芯片、终端、大模型、场景应用全链条的产业枢纽，能够赋能硬件厂商、AI 企业、终端消费者三方共赢。中国联通愿与产业伙伴携手，共同推动消费电子 AI 终端产业高质量发展，迎接全新智能终端时代。

来源：搜狐网

https://news.sohu.com/a/1052171775_116443

（二）中国工程院院士、鹏城实验室主任高文，德国国家工程院院士葛兴福（Ömer Sahin Ganiyusufoglu）作主旨报告

演讲核心观点：

主旨报告与主题演讲环节，院士、行业专家和企业代表围绕 AI 终端产业变革趋势、云边端协同创新、端侧大模型落地与生态共建路径深入研讨，从前沿技术、标准构建、产业落地、商业新模式多维度，为终端智能化高质量发展贡献前瞻思路与实践方案。中国工程院院士、鹏城实验室主任高文、德国国家工程院院士葛兴福，分别以《数字视网膜感知技术体系及其应

用》《个性化智能制造电脑、手机及智能眼镜》为题发表主旨演讲。

来源：搜狐网

https://news.sohu.com/a/1052171775_116443

十三、中国移动企业人工智能高质量发展论坛

（一）“智能驱动焕新 开放共创未来”——中国移动落实“人工智能+”行动创新发展实践(中国移动副总经理张冬)

演讲核心观点：

第一，以智筑基，让智能供给安全高效。中国移动全力构筑安全可信、高效普惠的新一代人工智能基座。一是加大 AI Infra 布局，推进全国一体化算力网建设，投产多个超万卡智算集群、12 个区域智算中心，升级 1500 个边缘智算节点，智算规模超 107EFLOPS。二是加快 AI Model 跃升，自主攻关九天全系列安全可信基础大模型，带动自主生态软硬件一体发展。三是加强 AIOS 锻造，升级移动模型服务平台 MoMA，建成规模化“词元超市”，让用户在高可信环境中选得放心、用得安心。下一步，将规模建设 AIDC，布局吉瓦级园区，攻关超大规模国产智算集群核心技术，构建多元异构算力生态，依托算网大脑实现算网电一体调度；突破九天安全可信大模型“根技术”，同时敏捷迭代 MoMA 平台，并加快建设智能体互联网与数联网。

第二，以智赋能，让智能扎根实体经济。中国移动打造开放的智能服务体系，实现全业可用、全民可享。一是支撑千行百业，研发 50+款行业大模型，落地 5000+个“AI+”项目，联合央企及行业龙头开拓高价值场景。二是服务千家万户，推动全量产品和服务嵌入 AI，焕新升级“灵犀”智能体，为 10 亿手机客户和 3 亿家庭客户提供智能生活服务；三是深入千城万村，打造“AI+”城市超脑平台，支撑构建“一屏观全域、一网管全城”的治理体系。下一步，将重点建设好制造领域汽车方向、农业领域农作物种植方向的两大国家人工智能应用中试基地；大力实施工业互联网创新发展工程，大会期间全新亮相工业互联网品牌“移动天工”，提供以移动云为基座、“天工”工业互联网平台为中枢的工业智能服务矩阵。

第三，以智聚势，让智能释放活力价值。中国移动持续深化产学研用联合创新。一是打造国家级 AI 开源创新载体，全力支撑国资央企人工智能开源“焕新社区”的建设运营；二是打造 10 个大模型产业创新基地，打造一批示范性强、可复制推广的 AI 标杆应用。下一步，将强化产业合作，贯通“算力—模型—应用”全链路；强化产学研投协同，打通“研-转-用”闭环；强化能力出海，聚焦“一带一路”共建国家，探索推进更多出海项目落地，进一步提升央企 AI 发展全球影响力。

来源：中国移动

https://www.10086.cn/aboutus/news/groupnews/index_detail_55490.html

(二) 上海东航数字科技有限公司及上海宝信软件有限公司等进行分享

在生态交流环节，中国移动分别围绕“移动天工”工业互联网产品体系、面向行业深度赋能的九天安全可信多模态智能底座、感通算智融合的 5G-A×AI 关键技术开展主题分享，上海东航数字科技有限公司及上海宝信软件(17.760, -0.83, -4.46%)有限公司分别就智慧民航、人工智能赋能重点行业的工业互联网平台进行分享交流。

来源：中国移动

https://www.10086.cn/aboutus/news/groupnews/index_detail_55490.html

十四、中国联通“模数同频 智造共振”AI 赋能新型工业化发展论坛

(一) 中国联通董事长董昕出席论坛致辞

演讲核心观点：

中国联通董事长董昕在致辞中表示，习近平总书记在 2026 世界人工智能大会暨人工智能全球治理高级别会议开幕式上发表主旨讲话，为我们进一步指明了方向、提供了遵循。中国联通认真落实“人工智能+”与“模数共振”行动，深化 5G、AI 与工业互联网融合，助力制造业高端化、智能化、绿色化发展，

将当好算网底座的供给者、应用场景的实践者与产业生态的营造者，与各方携手，更好地以人工智能赋能新型工业化。

一是筑牢先进的算网底座。落实“六张网”部署，构建“Agent+Token+AI 云”新模式，建强新一代通信网和算力网。在上海，中国联通推动 5G 专网深入制造现场，建设运营临港万卡智算中心，为工业 Token 创造、转运、存放、应用提供坚实保障。二是共拓鲜活的应用场景。依托联通格物工业互联网平台和元景 MaaS 平台，纳管超 1400 万台设备，沉淀超 100TB 工业数据集，培育超 75 个工业大模型，上线 150 多个工业智能体，在上海联合中国商飞、上海电气等落地一批标杆项目，将持续强化“数据供给-模型迭代-场景赋能”循环，探索具身智能、群体智能等前沿领域，让 Token 为企业创造更大价值。三是构建友好的产业生态。充分发挥联通 Universe 生态开放平台作用，依托“数据底仓、联通优选、解决方案工厂”三大能力，共建合作生态，共享 Token 发展红利。

董昕指出，中国联通将全面落实以习近平同志为核心的党中央的各项要求，把握守正创新、行稳致远发展主基调，聚焦“连接”“算力”“服务”“安全”四大核心赛道，打造更多“模数共振”示范标杆，助力加速新型工业化进程，为建设制造强国贡献更大力量。

来源：中国联通官方号

<https://baijiahao.baidu.com/s?id=1871241665160949193&wfr=spider&for=pc>

（二）元景筑原生，万悟启工智(联通数智副总经理丁鼎)

演讲核心观点：

联通数智副总经理丁鼎以《元景筑原生，万悟启工智》为主题展开精彩分享，系统介绍了联通元景工业智能体系的技术积淀、创新探索与产业落地实践，为工业领域 AI 规模化、可信化应用提供全新解决方案。

来源：工控网

<https://www.gongkong.com/news/202607/452016.html>

（三）AI 重构组织与运营——三一智能企业建设的新范式（三一集团 CIO 许国强）

演讲核心观点：

三一集团 CIO 许国强发表了《AI 重构组织与运营——三一智能企业建设的新范式》主题演讲，系统阐述了三一集团以 AI 驱动企业智能化转型的核心洞察与实践路径。他谈到：三一“全球化、数智化、电动化”的“三化”战略，构成了企业数智转型升级的顶层设计。

许国强重点介绍了“域智能体+数字员工”的落地模式，制造域智能体能力覆盖排产、调度、质量管控全链条，是“懂生产、稳交付、提质量、降成本的制造智能大脑”；生动讲述了三一与中国联通合作的典型应用实践“下料件自动排产”及其

成果：目前人工排产工作量锐减 60%，钢板利用率提升 1%，每年节约成本超千万元。他指出，AI 早已不是锦上添花的工具，而是重新定义组织与运营模式的底层驱动力。

来源：中国日报

<https://ex.chinadaily.com.cn/exchange/partners/82/rss/channel/cn/columns/sn19a7/stories/WS6a5d99dba310d709c2fbe6db.html>

十五、中国铁塔第二届“智擎未来”人工智能应用创新论坛

（一）中国铁塔副总经理刘国锋出席论坛致辞

演讲核心观点：

中国铁塔副总经理刘国锋表示，中国铁塔立足服务人工智能产业发展，依托独有资源禀赋，持续打造“位置+计算+电力+安全”四大核心能力，加速从“通信基建”向“数字基建”转型升级，聚力构建适配人工智能时代的新型数字基础设施。

刘国锋指出，人工智能的蓬勃发展，离不开坚实的数字基础设施作为支撑，数字基建作为 AI 产业发展的核心底座、智能时代的算力基石，是人工智能技术从理论研发走向规模化应用、从单点突破走向全域普及的根本保障，为人工智能场景落地、技术迭代、生态繁荣提供了全方位、全链条的坚实支撑。作为数字基础设施建设运营的“国家队”“主力军”，中国铁塔立足服务人工智能产业发展，依托独有资源禀赋，持续打造“位置+计算+电力+安全”四大核心能力。

来源：东方财富网

<https://finance.eastmoney.com/a/202607183811677548.html>

（二）其他行业专家分享

演讲核心观点：

论坛期间，来自中国信息通信研究院、奇安信、浪潮通信、亚信科技、中国电信天翼云等单位的行业专家，分别围绕企业

AI 深化应用、AI 安全运营、算网协同发展、Physical AI 落地、Token 运营等主题发表主旨分享，并在互动交流环节与现场嘉宾展开深度研讨，为 AI 产业融合发展建言献策。

来源：中国铁塔

<https://www.chinatower.com/Index/show/catid/17/id/1773.html>

华信咨询设计研究院整理

十六、智启具身论坛 2026

（一）模型与数据飞轮：冲破物理墙，驱动物理 AI 智能涌现（觅蜂科技董事长兼 CEO 姚卯青）

演讲核心观点：

姚卯青认为，理解物理 AI，首先要回到智能的两条主线：表征与闭环。表征，是对世界的抽象压缩；闭环，是智能体与环境之间持续发生的“感知、理解、决策、行动、反馈”循环。数字 AI 的成功，本质上是知识表征在数字环境中的成功；而物理 AI 的难点在于，它必须同时打通经验表征与真实世界闭环。

在姚卯青看来，物理智能要实现规模化，必须突破三道墙：一是数据墙，真实交互数据极其稀缺且获取成本高；二是表征墙，跨任务、跨场景、跨本体的统一物理表征尚未形成；三是闭环墙，真实试错代价昂贵、反馈缓慢，难以规模化。他表示，物理 AI 从“能演示”走向“能干活”，关键不在于单点模型能力，而在于能否形成可泛化、可优化、可进化的通用具身智能系统。

围绕这一目标，智元构建了“预训练—后训练—持续学习”的三阶段训练架构，沿着 VLA 与 WAM 两条路线持续演进，并推动二者走向统一的世界推理动作大模型（WRAM，World Reasoning Action Model），以此推出自研 GO-2 基座模型、Act2Goal 世界模型、GE-Sim 2.0，搭配 SOP、Genie Evolver

百台级分布式真机强化学习体系，从算法、仿真、真机训练层面打通规模化迭代闭环。

在数据层面，姚卯青介绍觅蜂科技打造了全链路物理 AI 数据基础设施，依托 MEgo 系列采集终端、MEgo Engine 治理平台与真机数据集 AGIBOT WORLD，构建分层数据供给金字塔与部署回流数据飞轮，从源头解决行业数据稀缺难题。姚卯青强调，当执行、回流、挖掘、更新形成自动化闭环，模型能力才能在真实世界中持续迭代。

这一技术飞轮已经在真实场景中获得验证。在江西南昌，传统 3C 平板量产产线完成规模化落地验证，产线作业成功率达 99.99%。“模型决定起点，数据定义终局。当飞轮转动，智能涌现只是时间问题。”姚卯青表示，智元、觅蜂希望与全球学术界、产业界和开发者社区共同推动物理 AI 在真实世界中加速演进。

来源：新浪财经

<https://baijiahao.baidu.com/s?id=1871290176113476527&wfr=spider&for=pc>

（二）数据驱动的具身基座模型（清华大学计算机系副研究员苏航）

演讲核心观点：

清华大学计算机系副研究员苏航围绕《数据驱动的具身基座模型》指出，具身智能正在经历一次根本性的预训练范式演

化：从面向单一机器人和特定任务的行为模仿，转向面向开放世界的通用能力预训练。未来具身基座模型的目标，不再只是学习“看到什么、如何动作”，而是形成对环境状态、物理规律、动作后果和任务进程的统一理解。

他认为，这一范式将沿着三条主线加速演进：数据从单一本体走向多源混合，模型从行为模仿走向世界建模，学习从离线训练走向真实部署中的持续进化。机器人轨迹、遥操作数据、第一视角视频和人类行为视频将被纳入统一的预训练体系，不同机械臂、人形机器人和应用场景中的经验，也将从彼此割裂的数据，转化为可迁移、可复用的通用能力。

苏航强调，具身预训练的核心逻辑正在从“看得更多、模仿得更像”，升级为“理解世界、预测后果、持续进化”。在这一趋势下，未来行业竞争的重点也将从单项任务成功率和演示效果，转向基座模型的通用性、迁移性和持续学习能力。只有完成从被动世界理解、统一世界建模到闭环自我改进的范式跃迁，具身智能才可能真正从“预编程的任务执行器”走向能够适应开放环境的通用智能体。

来源：新浪财经

<https://baijiahao.baidu.com/s?id=1871290176113476527&wfr=spider&for=pc>

（三）规模化真实世界数据（Physical Intelligence 研究科学家任至意）

演讲核心观点：

Physical Intelligence 研究科学家任至意则从规模化真实世界数据出发，展示了机器人策略模型如何在多任务、多场景训练中形成跨本体泛化能力。例如最新模型可将在 ARX 机械臂上学习到的衣物折叠能力，迁移至 UR5 双臂机器人完成同类操作，大幅降低新本体部署适配成本。机器人可处理未见过的橱柜、卧室整理、流体容器操作等长尾场景，验证多元真实交互数据对通用智能的核心驱动价值。

任至意介绍，最新 $\pi 0.7$ 模型实现突破性跨本体涌现迁移能力：仅在 ARX 机械臂训练折叠衣物数据，即可零适配在 UR5 双臂机器人完成同款精细操作，无需针对新本体重新微调，大幅降低机器人部署适配成本。

来源：新浪财经

<https://baijiahao.baidu.com/s?id=1871290176113476527&wfr=spider&for=pc>

（四）面向通用灵巧操作的人类数据与仿真扩展(Genesis AI 联合创始人王尊玄)

演讲核心观点：

Genesis AI 联合创始人王尊玄带来《面向通用灵巧操作的人类数据与仿真扩展》，分享了 Genesis AI 面向机器人灵巧操作的全栈方法：以 1:1 人类手部数据策略为基础，结合 Genesis World 仿真环境、多模态大模型 GENE 与部署数据飞轮，推动机器人在数据、仿真、模型和真实部署之间持续迭代。

王尊玄指出，通用灵巧操作的突破并不只是模型架构问题，更依赖高质量数据的规模化供给。随着预训练数据规模提升，模型在任务级微调后的效能和效率也有所提升。Genesis AI 通过接近人类尺度的机器人手与非侵入式手套数据采集方式，缩小人手与机器人手之间的数据鸿沟；同时，Genesis World 仿真环境不仅用于强化学习后训练，也用于模型评测，以加速模型迭代。

来源：新浪财经

<https://baijiahao.baidu.com/s?id=1871290176113476527&wfr=spider&for=pc>

（五）“研究与部署飞轮”战略（Dyna Robotics 联合创始人兼首席科学家马也骋）

演讲核心观点：

Dyna Robotics 联合创始人兼首席科学家马也骋分享了“研究与部署飞轮”战略，通过前沿研究、产品部署和真实世界数据反哺，形成技术迭代的自我强化循环。为支撑该飞轮，Dyna 构建了“预训练数据金字塔”，整合离线数据、在线机器人数据和

高质量部署数据，已积累超 20 万小时预训练数据。模型架构上，采用低频“推理模型”与高频“世界 - 行动模型”协同工作，兼顾高层认知与底层执行效率。

马也骋介绍，Dyna Robotics 商用基座模型 DYNA-1 在餐巾折叠这一高难度任务中，成功率高达 99.4%，并能在 24 小时无人干预下自主折叠超过 850 个餐巾，实现了从“实验室演示”到“工业级可靠性”的跨越。同时，该模型展现出卓越的数据效率，适应新任务仅需不到 1 小时的后训练数据。

可以看到，几位嘉宾的分享共同指向一个判断：物理 AI 的通用化落地，不是单一模型能力的突破，而是数据、表征、仿真、后训练与真实场景反馈共同形成的系统性进化。通用能力来自多源经验的规模化复用，可靠性则来自真实闭环中的持续迭代。

十七、中兴通讯全栈智算生态论坛

（一）中兴通讯执行副总裁谢峻石出席论坛并发言

演讲核心观点：

中兴通讯执行副总裁谢峻石表示，中兴通讯始终坚持“把复杂留给自己，把简单留给他人”的理念，坚定战略投入，深耕底层技术，持续完善全栈全域 AI 能力体系，打造极致协同、TCO 最优的端到端智算解决方案，共建开放共赢的 AI 生态，让合作伙伴更安心、行业客户更省心、终端用户更舒心，最终实现“AI 让生活更简单”。随后，论坛现场重磅发布中兴 OEX 超节点产品，其以开放解耦打通生态壁垒，依托底层协议标准化与接口统一，全面兼容国内多元 GPU 生态，支持客户结合业务特征灵活匹配算力。以架构创新重构智算底座，首创 OEX 正交无背板架构，实现二零线缆设计，内部互联成本大幅下降，叠加“OEX 算力集装箱”与前置式开发，有效压缩上市周期，确保“上线即领先”并显著降低试错风险；以多芯协同释放集群潜能，支持自主选择最优芯片组合，通过高速互联总线协议与大容量交换芯片，打通 GPU 机内高速互联与机间无损交换，充分释放“1+1>2”的集群效能；以极致性能拥抱 Token 经济，围绕吞吐、能效与体验做系统级协同优化，从芯片、算法、整机、集群到软件全栈拉通，确保每一瓦电力转化为最大 Token 产出，打造全生命周期成本最优的“AI 工厂”。

来源：C114

<https://www.c114.com.cn/news/127/a1314380.html>

（二）中国移动研究院副院长邵春菊出席论坛并发言

演讲核心观点：

中国移动研究院副院长邵春菊指出，聚焦 AI 时代 Token 经济产业发展的新范式，中国移动系统布局“三网三业”战略规划，打造 Token 生产高性能工厂，构建 Token 分发超级入口，全面提升智能服务水平，共育智能经济新形态；同时，她分享了中国移动研究院在建设业界领先的、覆盖从算力基础设施到大模型及智能服务的全栈人工智能生态评测体系，支撑中国移动智能服务提质升级方面的实践；同时，她面向全产业链发起合作倡议，呼吁各方协同联动、聚力共生，携手打造开放多元、共建共赢的 Token 经济产业新未来。

来源：C114

<https://www.c114.com.cn/news/127/a1314380.html>

（三）其他嘉宾发言（清华大学教授郑纬民、中兴通讯高级副总裁张万春、上海仪电所属智算科技公司国产适配中心主任牛红星、阶跃星辰联合创始人、CTO 朱亦博）

在主题发言环节，与会嘉宾围绕未来 AI 算力竞争展开充分交流。

清华大学教授郑纬民指出，智能体时代核心是低成本高等级 Token 产出能力，异构 PD 分离将推动推理系统分布式、异构化发展。

中兴通讯高级副总裁张万春表示，作为数字经济的筑路者，中兴通讯不追求封闭的领先，而是致力于生态的繁荣。中兴通讯将持续交付确定性高、TCO 最优的智算方案，携手合作伙伴，在 Token 经济的浪潮中开创智算未来的新篇章。

上海仪电所属智算科技公司国产适配中心主任牛红星表示，将继续秉持开放合作、生态共生的理念，携手产业链上下游伙伴深化技术协同创新，持续提升智算云服务能级。在全场景国产算力服务方面构建了 4 大核心能力，包括超大规模集群系统工程能力，弹性灵活智算云平台服务能力，自主可控国产生态链接能力，垂类应用场景支撑能力等。他还就仪电智算云 yicloud，国产适配能力，智爱赛思 OpenAI4S 社区等进行深入的探讨和分享。

阶跃星辰联合创始人、CTO 朱亦博指出，当 AI 从生成答案走向持续行动，产业竞争也将从单点模型能力，转向模型、系统、芯片与终端的全栈协同。阶跃星辰正围绕这一趋势持续布局，推动智能从云端能力走向真实世界，成为可规模化的人机协同基础。

来源：C114

<https://www.c114.com.cn/news/127/a1314380.html>

十八、基座大模型架构创新与生态合作论坛

（一）什么才是真正好的人工智能产品（商汤科技董事长兼 CEO 徐立）

演讲核心观点：

当系统级交付能力进入产品，它带来的不只是效率提升，而是生产单元的重新定义。那么，什么才是真正好的人工智能产品？

商汤科技董事长兼 CEO 徐立在演讲中表示：“好的 AI 产品，既扩展人的能力，更促进人的成长。不是替代，而是能力拓展，让原来没有能力的人获得能力，连接到更复杂的工作和表达。不是依赖，是个人成长，不只是替人做事，而是让人因为使用 AI 而变得更好。”小浣熊 Raccoon Work：把产业级 AI 生产力，带给每一个人

商汤小浣熊正在把 AI 能力扩展带入真实用户和真实任务，把产业级 AI 生产力带给每一个人。目前，小浣熊已服务超过 2000 万用户，周活跃用户达 300 万，近期增长达到 10 倍量级。数据背后更值得关注的，是用户正在把越来越复杂的任务交给 AI——从单次问答，进入研究、分析、文档、PPT 和完整 workflow。在小浣熊这些用户中，我们看到了 AI 在不同维度上延伸了人的能力：小浣熊让经验持续增值，让专业能力规模化。拥有 20 年施工审图经验的景观设计总工李博，借助小浣熊把自己的判断标准、审图流程和经验沉淀为 AI 能力，让一个人的

专业经验可以被结构化、复用和规模化，此前很多行业经验只真正掌握在少数人手中。小浣熊让兴趣连接世界，让个人创造成为新供给。历史爱好者崔城希望把 500 个人物的一生、人物关系和地理迁徙放进一张故事地图里，比如我们可以在这张地图里模拟王安石的一生——他的诗、旅行轨迹以及各方面成长。小浣熊帮助这位历史爱好者把兴趣转化为一个完整、可交付的作品。小浣熊让科技普惠，从工具无障碍走向能力无障碍。特殊教育领域的探索者曹文伟把小浣熊称为自己的“第二双眼睛”。作为视障者，他利用小浣熊以及语音交互的方式，打造出一套更加适合视障者的教学材料。真正的无障碍不只是让软件按钮更容易操作，而是让特殊人群获得过去无法独立获得的信息和能力。这些能力延伸的背后，是商汤过去十年持续把 AI 带入真实产业的积累。商汤理解真实世界，沉淀多模态认知；理解产业场景，积累业务知识；进一步进入业务流程，最终交付业务结果。徐立认为，AI 进入产业，不只是模型能力，更是业务理解、任务拆解、知识沉淀、 workflow 构建、质量评估和安全交付的综合能力。过去，这些能力依赖专家团队和项目制交付；今天，商汤希望对它们进行 AI 原生重构，形成一个人、创业团队和企业都能够调用的生产力底座。对于个人，小浣熊让一个人拥有更完整的能力，成为更高效的生产单元。对于 AI 原生初创组织，商汤参与共建《OPC 九级人才能力标准》，并发起国内首个 OPC 能力挑战赛，共有 65 万人参与，

形成 2 万名 OPC Learner 和 24 位 OPC Builder。OPC 代表的是一种新的组织形态：当个人能够调度多个 Agent、构建完整 workflow，一个人或一个很小的团队，也可以拥有过去只有完整公司才能具备的生产能力。对于企业，小浣熊已服务 8000+ 企业用户，并进一步形成教育版、政数版和金融版等行业产品——它们不是互相割裂的多个产品，而是同一套产业级 AI 能力在不同岗位、不同数据和不同业务规则中的具体形态。产业级 AI 能力正在重塑每一个生产单元，小浣熊让个人拥有更完整的能力，让创业团队以更低成本进入市场，也让传统企业完成更深层次的智能化升级。小浣熊最终希望做到的，是把产业级 AI 生产力，带给每一个人。Seko：把 Studio 级创作能力带给每一位创作者

在内容创作领域，商汤在做另一件事。过去，大多数人内容是内容消费者。能够制作电影、短剧、动画和专业视频的，往往是少数专业团队。商汤 Seko 希望让普通人从“看内容”，走向“做内容”。上线一周年，这个目标正在变成现实。从 1.0 行业首发短片创作智能体，到 2.0 实现百集连续创作，Seko 上线一周年之际，用户规模已突破 100 万，用户量增长 26 倍；企业客户达 1500 家，单月营收增长 89 倍；付费用户日均创作时长超 1 小时，用户增长与商业化质量同步验证，不断定义“AI 如何参与创作”。全新的 Seko 3.0 以“Agent 框架+超创 Skill 生态”为底座，打通了多类型意图理解、任务自动拆解、模型智能匹配和

workflow 组装的全链路闭环，实现了全品类内容的创作任务和自由表达，定义“视频创作领域的 Codex 时刻”。无论是短漫剧，还是音乐 MV、剧情种草等泛娱乐内容，亦或是直播间互动、科普讲解，Seko 3.0 都能实现角色、剧情、镜头的连贯稳定，大幅降低成片返工率。Seko 让不同类型的人，都能找到适合自己的表达方式。这些作品背后，是非常不同的人。有高校教师，用视频把知识变成孩子愿意看的内容；有小学生，用 AI 创作动画；有大学生制作系列短片，账号播放量突破 4 亿；有公务员的作品入选北京国际电影节 AIGC 单元；也有青年创作者从一个想法成长为创作者品牌。一位媒体记者“班壳”将北京春日花粉过敏的真实生活体验写成原创歌曲《花粉》，并借助商汤 Seko AI 完成全片音乐 MV 制作，她逐步把兴趣发展为自媒体副业，相关创作最终登上湖南卫视舞台。

再回到一个更基本的问题：什么是好的 AI 产品？徐立认为至少有两个维度：第一，不是替代，而是能力拓界——让原来没有完整能力的人获得能力，连接到更复杂的工作和表达。第二，不是依赖，而是个人成长——AI 不应该让人离开工具之后什么都没有留下，好的 AI 产品应该帮助人形成更好的习惯、建立更好的判断。商汤咔皮家族作为 Personal Intelligence（个人智能）产品矩阵，正是以 AI 赋能“个人成长”。咔皮健康、咔皮记账、Kapi 相机，分别作为 CWO·健康官、CFO·财务官、CCO·创意官，将 AI 化身每个人日常生活的“三位私人

高管”，它们共同帮助用户形成更好的习惯、做出更优的判断，成为更好的自己。截至目前，咔皮家族累计用户已突破 4000 万，年增速超 500%。小彩蛋~作为好的 AI 产品的现场检验，小浣熊在世界杯期间上线了足球分析 Skills，提前一个月将佛得角列为可能晋级的黑马，并在参与预测的 12 家大模型中排名靠前。

来源：商汤科技

https://mp.weixin.qq.com/s/8FnbU_q9LoVZadwXNIuVkQ

（二）“首席科学家”圆桌对话（商汤科技首席科学家林达华主持，复旦大学邱锡鹏教授，达摩院首席科学家赵德丽、阶跃星辰首席科学家张祥雨、新加坡南洋理工大学刘子纬教授）

核心演讲观点：

过去一个多月，GPT-5.6 等国际最新模型陆续亮相。林达华开场便指出，如今最先进的大模型几乎都已经把多模态理解能力变成“标配能力”。

但另一个问题随之浮现：多模态到底只是大语言模型能力边界的延伸，还是下一代智能真正的起点？进一步说，今天依靠 Next Token Prediction 建立起来的大模型，真的能够理解物理世界吗？

在高密度思想碰撞中，五位科学家围绕多模态的内生本质、核心瓶颈以及 5 年后的终极图景，展开了一场足以重构行业认知的“生死局”辩论。

01.

“走出洞穴”

再论多模态是“外挂”还是“灵魂”？

林达华一开场就抛出了这个核心问题，也是整场论坛最核心的一次思想碰撞。

当前业内存在两种截然不同的声音。一种认为，大语言模型已经成为未来 AI 的核心，多模态只是语言模型能力边界的自然延伸；另一种则认为，当 AI 真正从数字世界走向物理空间、开始构建世界模型时，多模态就是智能最基础、最原生的组成部分，而不是语言模型的外挂。

未来 AI 到底应当以语言为中心，还是以世界为中心？林达华把这个问题直接摆上了台面。

对此，邱锡鹏的观点鲜明：智能还是要依赖语言的。他认为，多模态未来的发展重点不是取代语言，而是如何把多模态信息与语言真正对齐。

不过，在他看来，今天的大模型距离真正理解世界还相差很远。“模型真正欠缺的，不只是多模态，而是 Context。”这里的 Context 是指真实世界复杂情境的理解能力。未来模型面对的不只是更多图片、更多视频，而是如何把现实世界复杂情境转化成模型能够理解的信息。未来多模态的发展，很可能需要一整套围绕现实环境构建的 Harness 以支撑模型理解真实世界。

与邱锡鹏的观点不同。赵德丽认为，基于多模态的模型，是下一代智能的范式。

他首先从更宏大的视角总结出通向智能的四条路径：大语言模型、走向开放空间的机器人、基于数字生命的模拟、基于神经信号的交互。这四条路径有一个共同特点，都与人类产生智能的源头信号有关，而且全都是多模态的，没有一个单一模态。

“大语言模型是知识的压缩，视频生成模型也是人类行为的压缩和记录。大语言模型可以学习智能，视频生成的模型一样可以学习智能。”赵德丽强调，他举了猎豹捕猎、蚂蚁寻路的例子，这些高超的技能都不是通过语言学习的，而是通过与环境的交互、模仿、反馈。

赵德丽更进一步指出，物理世界需要的是四维因果关系的建模：物理空间、物理属性、因果关系、动作与运动，而大语言模型只是一维的。因此，“无论从智能生发的原理还是算法的逻辑，基于多模态的模型，都是下一代智能的范式”。

与邱锡鹏和赵德丽的角度都不相同，张祥雨则把焦点转换了。

他认为核心不是语言还是视觉的路线之争，而是“学习范式”：怎么让智能体学会自主学习，或者叫后天学习、在线学习。“这个不解决，你不管是走语言还是视觉，其实都是一样的。”

为什么语言模型今天看起来更成功？张祥雨认为，我们现在的训练数据是静态的，是经过人类大量收集、清洗、整理得到的。语言天生信息密度高，在这种“模仿学习”范式下自然占优。而视觉信号稀疏、高维，在这套范式下工作得不好。

但问题来了，张祥雨话锋一转，当大模型走向智能体时代，一切都在变化。“智能体跟环境做交互，但并不能从反馈中高效、自动地学到东西，缺少了自我进化能力。”他举了一个生动的对比：“Codex 订阅一个人顶十个人，但用了多久提升都有限；招来的实习生一开始不行，通过实际工作很快就进步了。这是自我提升的能力，而我们的模型没有。”

眼看讨论在语言中心论和世界中心论之间形成了张力，刘子纬从哲学层面给出了一个富有深意的判断。他引用了柏拉图的洞穴比喻——语言是真实世界的低维投影，“我们在洞穴里观察投影，总结世界规律，但终归需要迈出洞穴，走向真实的世界”。

[OBJ]

具体到数据层面，他把语言比作化石燃料，“人类积累了很久，但终归有限，我们肯定会走向下一个时代，利用其他能源形态——多模态数据”。任务层面，“很多高价值的任务，比如制造业，一定脱离不了多模态，而且多模态真的能涌现出下一代智能”。

可以看到，从林达华的“语言中心与世界中心”之问，到邱锡鹏的“情境智能”、赵德丽的“第一性原理”，再到张祥雨的“持续进化”、刘子伟的“脱虚向实”，尽管各位学者的切入角度各不相同，但指向的趋势是高度一致的：多模态正在从“辅助工具”走向“核心基础设施”。

而这种转变的真正驱动力，是 AI 应用场景从数字空间向物理空间的不可逆迁移。

02.

数据、架构、范式

多模态下一次跃迁，缺的是什么？

从第一性原理展望长期 AI 前景之后，林达华把问题拉回现实：当前主流 AI 方法，比如数据、模型结构、训练范式，足以支撑我们走到下一阶段吗？真正的瓶颈在哪？

这个问题在业内已经有了明显的观点分野。一种判断是乐观的：当前的模型架构已经具备足够潜力，真正的限制在于高质量多模态数据仍然稀缺。特别是带有时间关系、空间关系、行为过程和结果反馈的数据，比互联网文本更难获取、标注和使用。只要解决数据与算力问题，现有路线仍有很大的提升空间。

另一种声音更为审慎：问题并不只是数据量不足，而是基于下一个 **token** 预测的训练方法，可能天然不擅长学习物理规律、因果关系和持续变化的环境。真正的突破可能需要原生多

模态架构、世界模型、强化学习、主动探索或持续学习等新的范式。

那么限制多模态 AI 继续突破的首要因素究竟是什么？扩大数据和算力是否仍然有效，还是我们已经接近必须改变模型架构和训练方式的阶段？

首先，刘子伟给出了一个冷静的判断：“这三个方面（数据、模型结构、训练范式）都不足以推向下一个阶段，都有改进的空间。”

他先从数据层面拆解了“为什么语言行、多模态不行”：语言模型撞上了好运，互联网 20 年积累的语料、为玩游戏而发展起来的 GPU 算力，所有要素恰好齐备。但多模态呢？“互联网上有很多视频图片，大部分是相对静态的，没有太强的长程关联。”做家务、每天工作这种长程任务很难被记录下来。“语言模型是无心插柳柳成荫积累数据的，多模态需要一个更好的范式来记录、保存、标注数据，这不仅仅是技术问题，更是组织问题、社会问题。”

紧接着，他把矛头指向了架构：“整体架构还是以语言模型为核心，多模态能力以桥接方式进入。”他提出了“原生多模态”的挑战：输入侧怎么做到原生？输出侧怎么做到原生？长上下文、稀疏编码、在哪一层做表征——这些问题都还没有答案。

最后落到学习范式，他承认这是“更大的问题”。“所有做视觉的人都有这个感觉，10年前大家就在讲这个问题，依然没有被解决。”

刘子纬的“三连击”把问题拆得很开，但张祥雨或许觉得还可以更彻底。他接过“持续进化”这个命题继续深挖。

他认为，自主学习是共识，但具体路径存在分歧。今天的模型有比较长的上下文，但“还远远不够长”，而且有严重的上下文退化。工程上靠上层 Harness 做上下文管理、做 Self Evolve，是一条有人在尝试的路线，但问题很明显。

多模态让这个问题变得更严重。“多模态的信息太长了，压缩多模态信号不像压缩文本那么自然。如果塞满了 100 万 token 的图呢？压缩率非常有限。”

然后他抛出了一个犀利的问题：“仅通过外部记忆，真的能让智能体本身产生演化吗？就像一个学生很会记笔记，但假如他的海马体无法再记忆新的东西，仅靠外部记忆，他能取得长远进步吗？不可以的。”

张祥雨继续拆解，下一代自主学习要解决三大问题：怎么学（在线学习问题）、学什么（探索与好奇心建模）、怎么学得高效（效率问题）。而当下最值得投入、最有可能短期内出成果的，就是在线学习。

在线学习又面临两个核心挑战，他继续展开，第一是层次化记忆建模，今天的 Transformer 建模短期记忆很优秀，但瞬

时记忆（感知记忆）完全没有，长程记忆又分为情境记忆和语义记忆，对应着不同的机制。“Transformer 的上下文从原理上更像一种工作记忆，它是无损的，但有效长度比较短。你在用一个试图把工作记忆变长的方式来建模原理上不同的记忆，这肯定有问题。”

第二是学习算法本身。很多人把赌注押在强化学习上，但张祥雨提出了不同观点：“RL 在过去推理时代很成功，数据效率很高，但这个效率很大程度来自 On Policy，自己生成数据再蒸馏回去。满足这个特性的不止 RL，蒸馏本身、Self Correction 都满足。”他认为 On Policy 可能是解决监督信号稀疏问题的出路，而世界模型则是解决效率问题的另一个关键。

林达华插进来追问了一句，把讨论引向 Memory 机制。“现在的模型，不管是 Dense 还是 MoE，本质上都靠 Context 承载记忆。Context 从 1 兆推到 2 兆，但有效的还是前面一小段。这是不是说明，我们用一个短期工作记忆的机制去承载本该属于长期结构化记忆的功能，需要有较大的革新？”

张祥雨立刻点头表示认同：“没错，这也是非常重要的一个前沿研究点：如何建模 Memory 机制。”人类的记忆机制是分层的，包括瞬时感知记忆、持续 2-4 秒的短期工作记忆、长期的情境记忆和语义记忆。Transformer 的 Context 从原理和特性上就更像工作记忆，它不会丢信息，但有效长度就是有限。当视觉信号输入以后，它就永远留在 Context 里，机制上不允

许压缩掉。而人类会主动选择、压缩、清除。这是根本性的差异。

此时，赵德丽加入了这场关于“瓶颈在哪”的讨论。他先提了一个行业现象：“模型架构逐渐收敛，算法本身没人关心了，大家接下来都去做应用”。作为算法出身的科学家，他认为这是巨大的误区。

“AI 从数字空间走向物理世界，首先面对的就是数据问题。”赵德丽指出，语言模型和视频生成模型的数据都来自互联网积累，而物理世界的数据极其稀少。“以机器人为例，去年谁能有 10 万小时数据就很牛了，今年大家才谈论大几十万小时，但远远不足以支撑 GPT-3.5 水平的训练。”

[OBJ]

▲几位科学家在同台探讨

不过，赵德丽话锋一转，给出了一个极具产业洞察的判断：“中国的工业场景、服务业场景、家电场景，多样性和需求远远大于欧美。所有数据都来自实际场景，从这个角度看，中国在具身方向上有天然的巨大优势。”

关于训练范式和算法架构，赵德丽强调了一个容易被忽视的事实：AI 有技术共享的传统，论文和技术报告公开，这“给人一个误区，认为技术不重要”。但实际上，“把强化学习引入基础模型，PPO、GRPO 这类具体算法，Sparse Attention、Sliding

Window 这些架构创新，DeepSeek 底层的优化机制——每一个技术进步都推动着领域发展”。

回到 Memory 创新，赵德丽认为它涉及两个层面：一是与芯片架构深度耦合的软硬件协同；二是个人数据和工作数据如何长期供模型使用。“DeepSeek 的 Ingram 算法把知识注入模型、耦合成一体、端到端训练，这是架构创新的一个具体范例。”

至此，关于“瓶颈是什么”的讨论呈现出四层递进：林达华首先定位问题，刘子伟拆解了数据、架构、范式三方面的局限，张祥雨深入记忆机制和在线学习的技术细节，赵德丽则从产业落地和数据获取的现实中找到突破口。科学家的视角彼此交错，又相互补充。

03.

高估 Agent，低估社会变迁：

五年预言

林达华把最后一个问题留给了时间尺度。回望 2022 年底-2023 年初的 AI 产业，现在发生的很多事完全预想不到。那么从未来 5 年后看现在，什么会被高估？什么会被低估？

林达华自己先分享了一段感受。2023 年初大家最焦虑的是 AI 会不会取代人类工作，但到今天看，AI 更多是在改变工作方式而非消灭工作，这可能是被高估的；而被低估的，是范式变革的速度，2023 年没几个人能预料到强化学习会在

2025-2026 年重新成为核心方法，也没人预料到“测试时扩展”会成为一个独立的竞赛维度。

邱锡鹏坦言：“预测一向是非常难的。未来 AI 大部分应该不是我们今天所想象的。如果刻意去追求某些东西，不一定会朝预期发展。相反，现在看起来不那么受关注的，反而可能得到充分发展。”

他给出一个判断：下一代的范式变革可能是“AI for AI”。“既然 Coding Model、Agentic Model 已经走出了突破，不如先把它走到极致，再让它去设计未来的多模态模型到底是什么样的，有没有更好的范式。”

邱锡鹏强调渐进式演化的力量。“ChatGPT 在 2023 年看起来是很大的范式革新，但它也是一点点走出来的。Transformer 还在其次，更前面的是 Seq2Seq 的范式，我可以统一所有自然语言处理任务。有了这种大的潜力认知，再去优化模型，直到某个突变的节点，突然成功了。未来肯定和我们现在预测的不太一样。”

与邱锡鹏的“无心插柳柳成荫”观点形成对比的是，赵德丽的预测相对具体。他认为，被高估的可能是“Agent 在具体企业级、商业级场景落地时的自主化能力”。被低估的则是“AI 对社会运作方式的改变”。5 年对 AI 来说是一个非常大的周期。

他举了一个很具体的例子：支付。“现在逐渐走入人与智能体、智能体与智能体之间的交互。这里涉及一个简单的问题，

智能体怎么做支付？整个支付体系都是围绕人构建的，移动支付是增加社会效率最有效的工具。但智能体支付呢？这背后是整个社会运作方式的重构。”

一个偏重“演化不可测”，一个偏重“社会确定性”。这两种预测方向看似迥异，实则互补：邱锡鹏提醒我们保持对技术路径的开放心态，赵德丽则提示 AI 将如何渗透社会肌理。

来源：智东西

<https://mp.weixin.qq.com/s/UbKwS-zBMF0WghStndMpRA>

十九、Enterprise AGI 价值落地与产业实证论坛

（一）没有企业专属数据，就没有企业专属智能（蚂蚁集团 CEO 韩歆毅）

演讲核心观点：

1、企业 AI 的智能不止于知识累积，更在于经验沉淀；企业 AI 的更高价值，不止于个体提效，更在于组织进化。真正懂企业的 AI 必须理解组织运转的隐性知识。2、依靠公开知识训练出的通用大模型，未必能够真正理解一家企业的运行逻辑。企业的核心能力往往藏在流程、决策习惯和文化基因中，而非公开文档里。3、实现企业 AI 深度落地的路径，是打造“组织智能”——没有企业专属数据，就没有企业专属智能。数据是企业 AI 竞争力的护城河。

来源：电子商务研究中心

<https://mp.weixin.qq.com/s/UeTqSNcdmwLs9oVj6T3oRA>

（二）原国家外经贸部副部长、中国入世首席谈判代表龙永图

演讲核心观点：

龙永图则从更宏观的层面，给这一轮 AI 浪潮提供了一个参照坐标。他将“入智”与“入世”作类比，认为人工智能对中国意味着新一轮新的制度改革、产业升级和开放机会。

就像当年放开外贸经营权催生了大批市场主体一样，AI 时代也将催生大量超级个体与“小而美”的小微企业，成为新

型市场主体。但他特别强调，技术的最终价值仍然要回到发展与民生：无论是“人工智能+制造”、还是“人工智能+消费”，都必须落实到更好的产业发展和更公平的民生改善上。

来源：蚂蚁数科

<https://mp.weixin.qq.com/s/T7FsPM31wkhUGws4u1tQOw>

华信咨询设计研究院整理

二十、共赢金砖论坛

（一）混合式 AI 是普惠的最优路径（联想集团副总裁陈敏仪）

演讲核心观点：

1、人工智能技术正在从共识走向实践，成为推动产业升级和社会经济发展的核心生产力。2026 年是 AI 从“热词”变“动词”的关键转折年。

2、终端侧的 AI 能力将成为 PC 和智能设备的新标配，端侧 AI 与云端 AI 的协同将决定用户体验的上限。

3、混合式 AI（端云协同）是 AI 普惠的最优路径——既保护用户隐私，又发挥云端大模型的强大能力。

来源：电子商务研究中心

<https://mp.weixin.qq.com/s/UeTqSNcdmwLs9oVj6T3oRA>

（二）让算力像水电一样随取随用（华为数据通信产品线副总裁王辉）

演讲核心观点：

1、提出“让每个 AI 设施紧密联接，资源共建共享，应用深度融合”的核心主张，AI 基础设施不应各自为政，而需形成统一的算力网络。

2、分享了华为 iMasterCloud 一站式 AI 基础设施生态解决方案，通过云网协同降低 AI 设施的部署与运维门槛。

3、AI 算力网络的未来在于“联接即服务”——让算力像水电一样随取随用，而非每家企业都自建算力孤岛。

来源：电子商务研究中心

<https://mp.weixin.qq.com/s/UeTqSNcdmwLs9oVj6T3oRA>

华信咨询设计研究院整理

二十一、其他论坛

(一) 终端是 AI 普及的“最后一公里”(腾讯副总裁林松涛)

演讲核心观点:

1、终端将是 AI 普及的“最后一公里”，终端形式不一定是 PC 或手机，甚至可能是具身智能。AI 要从云端走向用户身边，必须通过终端载体实现触达，这决定了 AI 普惠的广度与深度。

2、今年全球 AI 应用百花齐放，从以前“有问必答”的聊天助手类产品开始走向“使命必达”。AI 应用正在完成从信息提供者到任务执行者的关键跃迁，用户对 AI 的期待已从“告诉我”升级为“帮我做”。

3、真实场景永远是 AI 价值的第一现场，AI 要真正创造价值，必须先理解企业和个人面临的垂直领域需求。脱离具体场景的 AI 能力再强，也难以转化为实际生产力。4、模型更多决定 AI 上限，而工程决定下限——用户会因为一个好功能而使用，但能持续使用和信任，一定取决于工程化水平的下限。模型能力决定了 AI 能做什么，工程能力决定了 AI 能做多久、做多稳。

来源：电子商务研究中心

<https://mp.weixin.qq.com/s/UeTqSNcdmwLs9oVj6T3oRA>

（二）AI 智能体要在真实场景中“干起活来”（百度创始人李彦宏）

演讲核心观点：

1、中国拥有丰富的应用场景、完整的产业链和大规模的数字经济，海量的真实需求为智能体的落地提供了多元化的舞台。这是中国 AI 发展相较其他国家最大的禀赋优势。

2、AI 智能体有望走出实验室，在生产线上、在田间地头、在诊室里真正“干起活来”。智能体必须从虚拟空间走向物理世界，才能创造真实的经济价值。

3、AI 产业的下半场，比拼的不再是模型参数的大小，而是智能体在真实场景中解决实际问题的能力。落地能力决定胜负。

来源：电子商务研究中心

<https://mp.weixin.qq.com/s/UeTqSNcdmwLs9oVj6T3oRA>

（三）AI 能力正从“少数企业玩得起”走向普惠化（京东集团副总裁、京东云基础云业务负责人龚义成）

演讲核心观点：

1、京东正致力于建成全球最大的物理世界运营商，依托零售、物流、健康、工业等真实业务场景，推动 AI 向物理世界大规模落地。京东的优势在于拥有从供应链到消费端的完整物理场景闭环。

2、机器人“大脑”智能化程度偏低，核心原因之一是缺乏高质量的机器人物理世界数据。仿真数据无法替代真实物理环境中的交互反馈。

3、从去年到现在，京东的数据处理成本已降低 10 倍以上——高质量数据是前提，高性价比是数据走向规模化的基础。成本下降使得 AI 能力从“少数企业玩得起”走向“普惠化”。

来源：电子商务研究中心

<https://mp.weixin.qq.com/s/UeTqSNcdmwLs9oVj6T3oRA>

（四）医疗 AI 要回归智能化体系的本质（京东健康医疗 AI 负责人王国鑫）

演讲核心观点：

1、“京医千询 2.0”重点在拟人对话、可信推理和医学全模态三方面实现显著突破。“拟人对话”能模拟人类医生的问诊方式，进行多轮病情询问，按循证医学逻辑给出合理建议，而非简单的问答式交互。

2、“可信推理”是医疗 AI 的基石。“京医千询 2.0”严格遵循医生临床诊疗思维路径，注重内部知识与外部知识、思维深度与问题难度的结合，具备反思机制和专家反馈的自适应能力，并经超 140 个临床科室医生及百万级真实临床复杂病案的专家评测，确保医学推理可靠性。

3、“医学全模态”是大模型走向临床深水区的必由之路。
“京医千询”已实现对文本、影像、检验数据等多模态医学数据的综合解析，为精准诊疗提供全面支持。

4、京东健康强化“京医千询 2.0”不是要开发一个更大的模型，而是要回归医疗智能化体系的本质。“三引擎+四模型”架构是一项长期工作，将整体加快国内医疗智能化转型进程。

来源：电子商务研究中心

<https://mp.weixin.qq.com/s/UeTqSNcdmwLs9oVj6T3oRA>

（五）发展是第一位的，“不发展才是最大的不安全”（360创始人周鸿祎）

演讲核心观点：

1、人工智能正从大模型阶段迈向智能体阶段，AI 不再只是回答问题的工具，而是正在成为能够自主执行任务的“新物种”。AI 正在完成从“大脑”到“手”的进化。

2、对于 AI 产业而言，发展始终是第一位的，“不发展才是最大的不安全”。安全与发展并非零和博弈，不能用安全问题扼杀创新活力。

3、人工智能必须始终坚持安全、可信、可靠、可控的发展方向。没有安全边界的 AI，能力越强风险越大。

来源：电子商务研究中心

<https://mp.weixin.qq.com/s/UeTqSNcdmwLs9oVj6T3oRA>

（六）AI 要从“能想”到“能干”才算完成进化（网易有道高级副总裁刘韧磊）

演讲核心观点：

1、大模型解决的是“怎么想”，Agent 则解决的是“怎么干”——这种闭环交付能力正是用户真正需要的东西。从思考到行动，是 AI 价值跃升的关键一步。

2、当 Agent 真正走进学习、办公与生活的每一个场景，AI 才算完成了从“能想”到“能干”的进化。能力只有转化为行动，才产生真实价值。

3、教育 AI 的未来不在于提供更多答案，而在于引导学生提出更好的问题——培养批判性思维和创造力才是教育本质。

来源：电子商务研究中心

<https://mp.weixin.qq.com/s/UeTqSNcdmwLs9oVj6T3oRA>

（七）“技术越往前走，越要回到那个具体的个体”（智联招聘集团高级副总裁白岩）

演讲核心观点：

1、AI 对招聘求职生态的重塑，核心不是“变快了”，而是组织和个体之间的“连接关系变了”。AI 在极短时间内完成简历结构化处理和岗位匹配度排序，把 HR 从机械筛选中解放出来，省出的时间让 HR 回到更有价值的环节——与候选人深度沟通、判断潜力与气质，这些事 AI 做不了。

2、AI 带来的不是简单的岗位减少，而是岗位升级和重新分工。当下正在发生三件事：AI 能力正成为几乎所有岗位的通用门槛，市场、销售、产品都在要求 AI 协作能力；AI 相关岗位正从“单一行业”走向“全行业渗透”，航空航天、汽车、医药、新能源都在抢 AI 人才；企业端已从“利用 AI”转向“通过 AI 赋能人”。

3、未来 10 年，求职招聘的底层逻辑将从“筛选”变成“匹配”——企业和人才将基于能力、潜力、经历和组织需求做动态匹配。求职者端会出现“AI 代理人”，每个人都可能拥有长期陪伴自己的职业智能体，帮助理解趋势、管理能力、匹配机会。

4、组织形态将从科层制走向“人机单元”，效率提升从单纯依赖管理层级，转向人和 AI 协作带来的敏捷响应与持续创新。未来的组织不是“人+机器”，而是“人机融合的最小作战单元”。

5、智联的核心理念是让 AI 站在“人”这一边——帮人把自己的真实能力讲清楚，让组织看到更真实、更完整的人。“技术越往前走，越要回到那个具体的个体。”

来源：电子商务研究中心

<https://mp.weixin.qq.com/s/UeTqSNcdmwLs9oVj6T3oRA>

（八）2026 年是“世界模型元年”（昆仑万维董事长兼 CEO 方汉）

演讲核心观点：

1、2026 年是“世界模型元年”，标志着 AI 从内容生成走向对物理世界的理解与交互的关键转折。过去两年 AI 的核心命题是“生成”——生成文字、图像、视频、音乐；从 2026 年开始，核心命题将转向“理解”与“交互”——让 AI 真正理解物理世界的运行规律，并能与真实环境进行闭环互动。

2、三大产业预判：第一，世界模型将走出屏幕进入物理世界，具身智能正从实验室走向真实环境，要求机器人在行动前先进行环境推演——预判动作后果、评估环境变化、及时调整应对策略；第二，游戏行业将率先被世界模型改变，未来三到五年世界模型将成为游戏行业的基础设施；第三，AI 音乐、AI 视频将从技术试验变成大众创作工具。

来源：电子商务研究中心

<https://mp.weixin.qq.com/s/UeTqSNcdmwLs9oVj6T3oRA>

（九）每个传统行业都面临底层运行逻辑重构（华为公司董事、华为云 CEO 周跃峰）

演讲核心观点：

1、未来十年，智能体将重塑金融服务模式和决策体系。金融行业的高数据密度与强规则属性，使其成为智能体最先深度重构的行业之一。

2、过去依靠火电和水电稳定出力、源随荷动的传统运行逻辑，正在被“风光的间歇性”、“负荷的随机性”彻底重构，迫切需要以 AI 为核心驱动力重塑新型电力系统。

3、AI for Industries（产业 AI）正在从概念走向大规模落地，每个传统行业都面临一次基于智能体的底层运行逻辑重构。

来源：电子商务研究中心

<https://mp.weixin.qq.com/s/UeTqSNcdmwLs9oVj6T3oRA>

（十）AI 计算正从“堆卡时代”走向“优化时代”（无问芯穹联合创始人戴国浩）

演讲核心观点：

1、要避免智能计算基础设施出现功能繁多却非必需的问题，关键在于其在多元异构算力环境下的实际应用价值。算力基础设施必须“好用”而非“好看”。

2、瓶颈在于如何通过软硬件协同技术，将线性的技术供给匹配上非线性的需求增长。硬件的算力增长曲线与算法的需求曲线之间存在剪刀差，需要用软件效率弥合。

3、AI 计算正从“堆卡时代”走向“优化时代”——如何让每一张 GPU 发挥最大效能，比拥有多少张 GPU 更重要。

来源：电子商务研究中心

<https://mp.weixin.qq.com/s/UeTqSNcdmwLs9oVj6T3oRA>

（十一）青年一代要推动 AI 从“自动化”走向“自主化”（追觅科技创始人兼 CEO 俞浩）

演讲核心观点：

1、在“人工智能+”上升为国家战略的当下，青年科创力量既要攻克基础性科学难题，也要推动 AI 落地真实产业场景。青年一代是 AI 从论文走向产品的主力军。

2、青年一代应敢于挑战技术无人区，同时推动 AI 从“自动化”走向“自主化”，并让技术惠及全球。AI 的终极目标是服务全人类，而非少数国家或企业的专利。

3、青年科创者需在三个方向发力：以人才为根基培育中坚力量，以全球化为视野参与科技协作，以实业为底色坚持“技术+产业”融合。

来源：电子商务研究中心

<https://mp.weixin.qq.com/s/UeTqSNcdmwLs9oVj6T3oRA>

（十二）AI 辅助诊断帮助医生在复杂罕见病例中更快锁定方向（观壹智能联合创始人、CEO 谢伟迪）

演讲核心观点：

1、面向罕见病精准诊断，推出“DeepRare 2.0：基因-表型双链罕见病循证推理智能体诊断系统”，通过实践检验，该系统已在候选致病基因识别等行业评价中展现出专家级性能。

2、罕见病诊断是 AI 在医疗领域最具社会价值的应用场景之一——全球数亿罕见病患者长期面临“确诊难”，AI 正在缩小这一诊疗鸿沟。

3、AI 辅助诊断的核心价值不是取代医生，而是帮助医生在复杂罕见病例中更快锁定方向、减少漏诊。

来源：电子商务研究中心

<https://mp.weixin.qq.com/s/UeTqSNcdmwLs9oVj6T3oRA>

(十三)交易型 AI Agent 有望成为服务消费的新基础设施
(滴滴集团 AI 应用负责人刘佳奇)

演讲核心观点:

1、出行是连接消费者与线下消费场景的重要枢纽，平均每 1 元打车消费可带动数倍的其他线下消费。AI 正在放大出行作为“消费连接器”的经济乘数效应。

2、AI 小滴已拥有 90 多个标签，满足不同场景出行需求，从“用户适应产品规则”转向“产品理解用户需求”。AI 让出行服务从标准化走向千人千面。

3、交易型 AI Agent 有望成为服务消费的新基础设施。未来的 AI Agent 不仅能帮用户打车，还能基于出行场景主动推荐餐饮、娱乐等消费服务。

来源：电子商务研究中心

<https://mp.weixin.qq.com/s/UeTqSNcdmwLs9oVj6T3oRA>

(十四)安全应从设计之初融入整车架构(曹操出行首席科学顾问、图灵奖得主约瑟夫·希发基思)

演讲核心观点:

1、构造即安全——安全应从设计之初融入整车架构，而非事后打补丁。安全必须是自动驾驶的基因而非装饰。

2、系统协同——Robotaxi 是车辆、云平台、调度系统、运营体系组成的复杂系统，需协同运行实现整体可信。单一环节的安全不等于系统整体的安全。

3、运行保障——安全保障须贯穿运行全过程，通过实时监测、系统冗余、远程支持等机制维持安全边界。安全是动态的、持续的过程。

来源：电子商务研究中心

<https://mp.weixin.qq.com/s/UeTqSNcdmwLs9oVj6T3oRA>

（十五）“技术决定能不能上路，运营决定能跑多远”（如祺出行 COO 韩锋）

演讲核心观点：

1、“技术决定能不能上路，运营决定能跑多远。”Robotaxi 小规模试跑阶段重点是解决单车技术问题，但规模扩张到千辆乃至万辆级别，必须打通调度、充换电维保、远程云控、安全兜底与成本管控等全链条。运营能力是规模化的核心门槛。

2、Robotaxi 本质是出行服务，规模化落地的关键门槛在于运营能力，这正是出行平台的先天优势所在。技术公司擅长造车，出行平台擅长管车——各司其职才能加速落地。

3、Robotaxi 行业发展，政策和监管必须先行。如祺出行通过开放式 Robotaxi 运营科技平台，强化行业与政府监管之间的桥梁作用，为自动驾驶规模化、规范化发展提供可复制的实践路径。

4、数据共享的核心目标是驱动产业整体升级，而非仅服务单一企业技术迭代。两大关键方向：一是基础性、公共性数据共享与跨行业协同应用；二是企业高价值极端场景数据合规有序流通，通过可信数据空间、隐私计算等技术手段集中攻克行业共性瓶颈。

5、坚守“先锁安全底线，再释放数据价值”原则，严守数据采集、处理、应用各环节合规红线。通过常态化运营的Robotaxi与智驾数据采集车相结合，持续以低成本合规采集真实场景数据，构建“真实场景数据持续输入、合规高质量数据服务稳定输出”的数据飞轮。

来源：电子商务研究中心

<https://mp.weixin.qq.com/s/UeTqSNcdmwLs9oVj6T3oRA>