

加拿大-亚洲人工智能与数字创新理事会 (CAI) & 哈佛肯尼迪学院校友
多边数字创新与技术治理圆桌会议（第四期）· 政策白皮书 · 2026年第1号

从高层原则到跨国执行架构

前沿模型风险、主权冲突与协同治理路径

人工智能治理白皮书

联合起草与学术贡献：
CAI x HFTC 政策研究小组
加拿大-亚洲人工智能与数字创新理事会
(Canada-Asia Council for Artificial Intelligence and Digital Innovation, CAI)
哈佛肯尼迪学院 (Harvard Kennedy School, HKS) 校友联合倡议
& 哈佛前沿技术俱乐部 (Harvard Frontier Techne Club, HFTC, 暂定名)

2026年6月 · 联合发布

执行摘要 (Executive Summary)

在过去三年中，全球人工智能（AI）的讨论与政策工作高度集中，主要聚焦于模型能力、AI安全和顶层监管框架。然而，CAI x HFTC的研究表明，真正的战略脱节不在于政策文本的数量，而在于“治理与执行之间的缺口”（Governance-to-Execution Gap）。

国际社会已陆续推出了如《欧盟人工智能法案》（EU AI Act）、美国第14110号行政令（U.S. Executive Order 14110）以及中国《生成式人工智能服务管理暂行办法》等监管机制。这些制度性努力在风险分类、合规边界和责任划分方面发挥了至关重要的作用。然而，它们仍主要以国家主导的“自上而下”监管模式为中心，在应对前沿AI模型的快速迭代和系统复杂性时，日益显现出结构性滞后和执行僵化。

与此同时，全球AI治理正从单一法律框架向“多层协同治理架构”（multi-layered collaborative governance architecture）的形成初期过渡：

- **国际层面：** 联合国于2026年启动了“人工智能治理全球对话”（Global Dialogue on AI Governance），标志着193个成员国首次系统性地参与AI治理讨论，预示着AI治理在全球政治协商机制中的制度化。
- **国家政策实验层面：** 新加坡资讯通信媒体发展局（Infocomm Media Development Authority, IMDA）率先推出了针对“代理智能”（Agentic AI）的治理框架，首次将治理范围从“模型”扩展至“自主AI系统”，强调人类责任链和系统级风险控制。
- **国家战略层面：** 加拿大的国家人工智能战略（“AI for All”）作为核心参照点，强调研究创新与社会信任的融合。它为“中等强国”的AI治理提供了独特的基准，侧重于构建协同生态系统，而非纯粹的反应式监管。
- **技术标准层面：** 加州大学伯克利分校（Berkeley）等研究机构正在建立针对代理智能的风险评估与安全控制方法体系，包括工具使用风险、自主行为约束和系统可解释性评估，推动事实上的技术治理标准形成。

在此背景下，诸如具备高水平自主规划和跨任务执行能力的系统等前沿模型，正在打破传统“静态模型”的界限，在复杂环境中展现出显著的主动性和适应性。这对传统的、基于规则的监管方法提出了系统性适应挑战。

白皮书核心主张：未来十年的核心竞争力

未来十年，真正决定国家和地区核心竞争力的是“人工智能执行能力”（AI Execution Capacity）。CAI x HFTC将其定义为：一种核心系统工程能力，使一个国家或地区能够将前沿学术研究、制度治理政策、国际资本资源以及本地化的多元产业场景深度整合，并实现可扩展、安全且合规的部署。

为了弥合当前AI发展中研究、治理、资本与产业应用之间的结构性脱节，本白皮书提出了“CAI x HFTC枢纽框架”（CAI x HFTC Nexus Framework）。该框架旨在通过构建跨区域协作机制，提升不同功能节点之间的连接效率与制度协同。

该框架由四类功能节点组成：研究节点（Research Nodes）、治理节点（Governance Nodes）、资本节点（Capital Nodes）和部署节点（Deployment Nodes）。它们共同构成了面向复杂不确定性的“全球人工智能协调网络”（Global AI Coordination Network），为AI系统的大规模应用提供更具韧性的制度基础设施。

在此网络中，加拿大可以作为连接不同区域创新生态系统的“桥梁经济体”。凭借其在AI研究方面的长期积累、深厚的社会信任基础、多元文化结构以及七国集团（G7）成员国身份，加拿大具备促进跨区域AI合作与制度对接的独特潜力。

同时，该框架将与哈佛肯尼迪学院（HKS）政策研究社区的学术网络建立联系，包括哈佛前沿技术俱乐部（HFTC，暂定名）所推动的前沿技术与治理讨论机制，以提升在AI治理、技术演进和全球制度协调方面持续开展研究与对话的能力。

通过探索“加拿大-亚洲连接机制”（Canada – Asia Connectivity Mechanism），加拿大能更有效地将北美的技术创新与亚洲多元的资本和应用场景连接起来，从而在全球AI协调网络中发挥关键的协调与连接作用。

第一部分：CAI x 哈佛 HKS 对话亮点

在2026年6月举行的“CAI枢纽系列IV：多伦多-波士顿人工智能安全与治理圆桌会议”期间，来自哈佛肯尼迪学院（HKS）、伯克曼·克莱因互联网与社会中心（Berkman Klein Center for Internet & Society）、多伦多大学、加拿大全球事务部

（Global Affairs Canada）及业界的专家学者，就当前AI治理的核心挑战、风险特征及执行路径进行了深入探讨。基于现场发言的严谨记录，本报告梳理出以下五大核心洞察：

- **洞察1：AI治理正成为地缘政治的焦点。**

圆桌会议开场，哈佛肯尼迪学院 / 伯克曼·克莱因中心学者、世界知名安全技术专家布鲁斯·施奈尔（Bruce Schneier）引用了最新爆发的事件：美国政府因国家安全担忧，强硬要求前沿AI实验室Anthropic撤回其刚发布的高级模型Fable。这一里程碑事件表明，前沿AI模型的发布、审查与管控已不再仅仅是自主的商业决策或简单的技术伦理问题，而是上升到了国家主权意志博弈的焦点。前沿实验室在大国竞争的缝隙中面临前所未有的国家安全审查。与会者指出，当前的AI发展存在于大国主导的竞争结构中；这种地缘政治结构使得技术自主与政府监管之间的冲突成为核心治理议题。

- **洞察2：数据主权（Data Sovereignty）正成为国家AI战略的核心要素。**

加拿大全球事务部（Global Affairs Canada）高级政策分析师、哈佛肯尼迪学院MPA校友赛义夫·马哈茂德（Saief Mahmood）以个人学术身份指出，过去十年支撑大模型野蛮生长的“无界数据抓取”时代已经结束。主权国家正在迅速建立严格的数据边界。例如，加拿大在其最新的国家AI战略中，将“主权”提升到了前所未有的核心高度。加拿大日益倾向于将加拿大公民数据资产保留在国内，以满足模型训练或微调的需求，这反映了其自身及其他发展中国家的价值观。CAI x HFTC系列组织者、哈佛MPP、拉贾瓦利研究员（Rajawali Fellow）龚博雅（Selina Gong）从相反方向进行了论证，指出AI的数据访问边界是一个严重缺失的治理议题。我们必须精确定义AI不应访问的信息范围，特别是在医疗等高度敏感领域。学者们还指出，数据主权是“中等强国”（middle powers）和全球南方（Global South）国家在面对大型科技公司时保持话语权和议价能力的关键。

- **洞察3：AI赋能的系统性安全风险亟需关注。**

布鲁斯·施奈尔从宏观视角指出，AI作为一种“力量倍增技术”（force-multiplier technology），存在系统性风险（如能力不平衡带来的“锯齿状边界”以及不可预见的“阿拉丁神灯效应”），若缺乏国际协同防御机制，可能对全人类构成安全威胁。

多伦多大学坦特里医学院（Temerty Faculty of Medicine）助理教授、执业血液病理学家周强华（Steven Zhou）从生物医学与生物安全的交叉领域发出了新的警告。他指出，虽然AI加速了基因组分析和蛋白质结构解析，但若与基因编辑等技术结合，可能会产生不可逆的生物安全风险。他强调，此类风险不仅存在于大型实验室；个人在“后门”进行危险实验的管理难度更大。

• 洞察4：核心挑战是“执行架构”而非“原则”。

与会者达成共识：当今世界匮乏的不是“原则”，而是能够跟上AI高频迭代速度的“执行基础设施”（Execution Infrastructure）。布鲁斯·施奈尔提出了一个精辟的结论：“治理移动缓慢，而技术移动迅速。”传统的立法审查程序耗时数年，无法跟上AI在12个月内颠覆软件开发和企业生产力的技术节奏。未来治理的成功完全取决于“监管翻译”（Regulation Translation）的高效执行以及敏捷危机响应的韧性机制。

CAI x HFTC系列组织者、加拿大-亚洲人工智能与数字创新理事会（CAI）联合创始人、哈佛肯尼迪学院校友刁孝力（Stanley Diao）表示：“AI治理已不再是一个原则问题，而是一个跨司法管辖区的执行架构问题。”与会者认为，建立能够高效翻译技术迭代并敏捷响应的“执行基础设施”是未来治理成功的关键。

• 洞察5：多元参与者与全球南方的代表性不足。

赛义夫·马哈茂德直言不讳地指出，许多科技公司的盈利甚至超过了大多数发展中国家的GDP。这种权力不对称导致了像加拿大这样的“中等强国”和全球南方国家在模型设计、伦理标准和政策制定方面的影响力有限。即将入读哈佛肯尼迪学院MPP的Sky9 Capital投资分析师肖易菲（Ella Xiao）等与会者指出，当前的AI治理体系高度集中且呈寡头垄断态势，治理权力主要掌握在美国和中国的大型科技公司手中。与会者呼吁未来治理应基于双边或多边“志同道合”（like-minded）联盟，构建更具包容性的国际协调网络。

第二部分：全球AI治理文献综述

为了给白皮书注入深厚的制度理性，本节系统回顾了来自哈佛大学的最新文献与前沿概念框架：

“哈佛的独特贡献不仅仅是AI创新，更是AI时代的制度建设。”

1. 哈佛阿什民主治理与创新中心（Harvard Ash Center）：权力共享自由主义

正如哈佛大学丹妮尔·艾伦（Danielle Allen）教授及其团队在《治理AI路线图》（A Roadmap for Governing AI）中所述，当前的技术治理模式往往陷入“被动防御”和“惩罚性监管”的狭隘陷阱。哈佛团队创新性地提出了基于“权力共享自由主义”（power-sharing liberalism）的主动治理观。该框架强调，AI治理的终极目标应是促进“人类繁荣”（Human Flourishing）。这不仅需要保护消极自由（如隐私和防止歧视），更需要积极创造机会以促进积极自由——即保障社会各阶层广泛、深入参与“驾驭”技术发展方向的政治权利。哈佛路线图主张将权力从少数资本和科技垄断者手中重新分配给对公众负责、由民主引导的公共机构，同时在公共物品、人才培养和民主决策能力方面实施大规模的国家战略投资。

2. 哈佛副研究员工作论文250号：双重指令

哈佛高级研究员保罗·卡尔沃（Paulo Carvão）在其2025年的论文《双重指令：AI时代的创新与监管》（The Dual Imperative: Innovation and Regulation in the AI Era）中，深入分析了“加速主义者”（Accelerationists）与“末日论者”（Doomers）之间的极化现象。卡尔沃指出，两者都陷入了非黑即白的虚假二分法。他的“双重指令”（Dual Imperative）证明，前沿AI模型极度扩张带来的复杂不确定性和系统性风险，无法完全通过现有的边界修复或算法对齐来解决。因此，社会必须采取双向交织的路径：一方面，通过持续的技术发明跨越技术奇点以遏制灾难性后果；另一方面，必须依靠“智慧监管”（Smart Regulation）重塑私营部门的激励机制，强迫其将逐利动机置于公共责任之下。

3. 《哈佛法律评论》第138卷：协同治理框架

在《法律发展——人工智能》（Developments in the Law — Artificial Intelligence）第三章中，哈佛法律界聚焦于前沿AI生态系统的“开放-封闭光谱”（Open-Closed Spectrum）。文献指出，开源AI（如Meta的Llama系列）在本质上对“技术民主化”起到了关键作用。通过赋予公众关于模型架构和代码的四项基本自由（运行、研究、修改、分享），它极大地分散了技术霸权并激发了基层创新。然而，开源也带来了巨大的滥用风险，如低成本的“无审查模型”和安全“紧急停止开关”（kill-

switches) 的剥离。为了打破这一困境，该文提出了“协同治理”（Co-Governance）制度创新模式：彻底摒弃传统的一刀切、自上而下的监管路径，引入包括公民议会、定期审议机制和多利益相关方参与在内的持续参与模式。这一模式将监管前线推进到技术的整个生命周期，确保政策能随技术动态迭代而实施敏捷的自我修正。

第三部分：新兴趋势

综合前沿圆桌会议的观点碰撞与全球顶级文献脉络，CAI x HFTC梳理出未来5到10年全球AI生态在执行层面的五大不可逆前沿趋势：

趋势 维度	现状（原则层）	未来方向（执行层）
AI安全	专注于事后对齐（RLHF）、静态补丁修复及简单的商业合规审查。	演变为实时全生命周期监控，以及针对具有动态对抗和高主动性能力模型的国家安全主权“紧急停止”机制。
AI主权	盲目迷信无界数据全球化及单一的跨国巨头监管标准。	全面构建主权数据边界；主权大模型、本地化数据微调和本地化计算信任成为主权国家的硬性标准配置。
代理治理	将静态“黑盒”模型作为唯一监管对象；无法定义自主行为代理的法律责任。	转向行为控制，并针对具备自主问题解决和规避能力的“代理智能”（Agentic AI）进行实体穿透式的法律责任界定。
计算治理	将计算视为普通商业IT资源和硬件供应链，由市场价格配置驱动。	计算上升为战略性、稀缺性公共资源；代币税（Token Tax）、能耗控制和计算外交成为国家地缘竞争的核心工具。
全球南方包容	治理文本层面的道德宣示，实质性利益再分配被完全边缘化。	转向“AI发展红利再分配”、“大模型共有权”及跨国补偿性技术转移网络。

第四部分：CAI x HFTC 人工智能治理框架

基于上述学术洞察与趋势分析，本白皮书正式发布CAI x HFTC的核心原创智力产品：CAI x HFTC框架——连接研究、治理、资本与部署。

CAI x HFTC框架旨在彻底打破当前系统中“研究是研究，融资是融资，监管是监管，产业是产业”的四种结构性缺口（Gaps）：

- 缺口1（研究→政策）：学术研究与政策制定之间缺乏高效的翻译机制，导致前沿科学成果难以及时进入政策和监管设计流程，造成监管框架形成时与技术发展之间的时滞。
- 缺口2（政策→资本）：在某些情况下，政策与合规体系缺乏对创新风险的差异化识别能力，可能在减少系统性风险的同时，抑制了早期高风险但高潜力技术的投资，影响了资本在安全与前沿技术领域的精准配置。
- 缺口3（资本→部署）：资本进入技术应用阶段时，往往缺乏与具体产业场景的结构化对接机制，导致资金从技术研发向物理应用转化过程中的匹配效率不足，影响了关键领域的规模化落地。
- 缺口4（部署→全球扩张）：在跨区域扩张过程中，本地成功的技术应用需要适配不同的司法体系、监管标准和数据治理规则。这些多制度环境的差异构成了全球扩张的协调与合规挑战。

CAI x HFTC框架：五级协同架构

- 第1层：研究节点 (Research Nodes)

以哈佛大学、多伦多大学、向量研究院（Vector Institute）和斯坦福大学等全球研究机构为核心节点，为前沿模型研究和AI对齐提供理论支撑。

- 第2层：治理节点 (Governance Nodes)

联合政府机构、监管部门和国际组织，推动研究成果转化为政策框架和跨国监管协作机制，提升不同司法管辖区之间的互操作性。

- 第3层：资本节点 (Capital Nodes)

连接全球风险投资机构和长期资本配置方（包括家族办公室等），支持AI安全、可信度和基础设施相关技术的长期投资与可扩展发展。

- 第4层：部署节点 (Deployment Nodes)

推动AI系统在医疗、高端制造和能源系统等关键物理场景中的安全部署与渐进式应用验证。

- 第5层：全球协调节点 (Global Coordination Nodes)

作为跨区域连接与信息协调层，促进不同地区在标准、治理实践和风险响应机制方面的对话与协作，支持全球AI系统在多制度环境下的稳定运行。

第五部分：为何加拿大至关重要

在全球科技竞争的叙事中，绝大多数分析报告将目光锁定在美中权力矩阵的碰撞上，或聚焦于欧洲在立法层面的防御性繁文缛节。CAI x HFTC框架在此提出了一个独特的差异化观点：在全球AI执行网络中，加拿大是一个被严重低估的“中立AI强国”（Neutral AI Power），扮演着不可或缺的跨国协调枢纽角色。

加拿大拥有以图灵奖得主杰弗里·辛顿（Geoffrey Hinton）为核心的现代深度学习学术根基，并拥有包括向量研究院（Vector Institute）、Mila和Amii在内的全球顶级科学高地。同时，加拿大独特的不具攻击性的地缘政治定位，意味着其不带有技术霸权的色彩，具备全球稀缺的“中立性与社会信任”。作为七国集团（G7）的核心成员和亚太经济合作的重要力量，加拿大在北美政策圈、欧洲智库圈和亚洲金融体系中享有极高的多边政治与经济信任。

因此，加拿大在塑造AI治理与部署的未来方面，处于连接北美与亚洲的独特地位。通过这条中立管道，北美的前沿研究可以被无缝翻译为考量亚洲市场文化需求的各种安全执行系统，从而形成一条横跨太平洋的“超级信用走廊”。

第六部分：CAI x HFTC 2026-2030 年议程

未来五年，CAI x HFTC将致力于构建一个连接研究、治理、资本与产业实践的“全球人工智能执行网络”，推动AI从概念创新走向可信治理与规模化实施。以治理对话、产业部署、人才培养、资本协作和国际合作为五个核心方向，CAI将持续推动加拿大、亚洲及全球创新生态系统之间的交流与协同。

- **倡议1：全球AI治理对话系列与哈佛前沿技术俱乐部**

CAI将继续推进与哈佛及其他学术网络的AI治理对话机制，并支持由龚博雅创立，并在会议期间正式宣布的“哈佛前沿技术俱乐部”（Harvard Frontier Techne Club，暂定名）。该俱乐部由CAI Nexus合作参与发起，立足于哈佛社区，总体倾向于亚洲视角，对跨背景人士开放，关注AI及新兴技术与未来议题，强调对未来发展的预见性“演绎”探讨以及跨学科、跨区域的交流。

- **倡议2：AI部署实验室 (AI Deployment Labs)**

CAI x HFTC将联合高校、研究机构 and 行业伙伴，共同探索AI在医疗、能源系统、供应链和公共服务等领域的实际应用，推动可信AI（Trusted AI）从研究走向产业化落地。

- **倡议3：全球AI奖学金计划 (Global AI Fellowship)**

CAI x HFTC将培养一代具备技术、政策与国际视野的新型AI领导者，汇聚来自全球顶尖大学和机构的优秀青年人才，推动跨学科合作与国际交流。

- **倡议4：可信AI资本网络 (Trusted AI Capital Network)**

CAI x HFTC将连接北美与亚洲的投资机构、创新基金及行业资源，支持负责任AI与硬科技创新项目的发展，促进资本与创新生态系统之间的长期合作。

- **倡议5：加拿大-亚洲负责任AI创新门户 (Canada – Asia Responsible AI Innovation Gateway)**

CAI x HFTC将促进加拿大与亚洲在AI领域的人才交流、创新合作与政策对话，在尊重数字主权和监管要求的基础上，推动开放、可信且可持续的国际合作。

CAI x HFTC 原则

秉持负责任创新（Responsible Innovation）、可信部署（Trusted Deployment）、数字主权（Digital Sovereignty）、国际合作（International Cooperation）及公共利益（Public Benefit）的核心理念，CAI x HFTC致力于推动人工智能向着更安全、更包容、更具全球影响力的方向发展。

附录

I. 演讲者洞察总结

演讲者姓名	机构与身份	核心发言总结与原创政策观点
布鲁斯·施奈尔 (Bruce Schneier)	哈佛肯尼迪学院 / 伯克曼·克莱因中心研究员，安全技术专家	1. 治理错配：技术迭代快于治理。2. 反对政府参股：会产生巨大的利益冲突。3. 企业法律责任：AI公司必须为模型输出负责。4. 公共AI愿景：建立非垄断性的公共AI模型。5. 不可控后果：“阿拉丁神灯效应”——以不可想象的方式达成目标。6. 垄断是政策问题：反垄断疲软的结果。7. 代理风险：从聊天机器人转向智能代理改变了治理面貌。
赛义夫·马哈茂德 (Saief Mahmood)	哈佛肯尼迪学院MPA	1. 复杂性：治理AI是巨大挑战。2. 国家化：可能导致危险路径。3. 中等强国声音：虽非最大，但需要话语权。4. 数字主权本质：对数据、基础设施、采购的控制，而非自给自足。5. 治理碎片化：转向双边/三边“志同道合”伙伴关系。6. 三大支柱：信任、适应与主权。7. 生物风险：缺乏护栏的AI生物安全日益令人担忧。
周强华 (Steven Zhou)	多伦多大学坦特里医学院助理教授，血液病理学家	1. 医疗AI封闭系统：更可控。2. 基因武器风险：可能被用作武器。3. 不可逆演化：AI辅助基因编辑可能改变人类方向。4. 个人生物威胁：难以管控“后门”实验。
龚博雅 (Selina Gong)	CAI x HFTC组织者，哈佛MPP，拉贾瓦利研究员	1. 能力导向治理：缺乏预训练治理架构。2. AI代理/不稳定性：行为可能不一致。3. 数据边界：AI不应访问哪些信息？4. 跨国AI社区：发起“哈佛前沿技术俱乐部”。

刁孝力(Stanley Diao)	CAI x HFTC组织者, CAI Global联合创始人, 哈佛肯尼迪学院校友	1. 从原则到执行架构: 如何跨司法管辖区运作? 2. 系统性风险: 风险变得结构化且跨系统。3. 国家能力设计: 聚焦计算与数据基础设施控制。4. 范式融合: 在实验室、学术洞察与治理之间。5. 参与发起发起“哈佛前沿技术俱乐部”。
Angel Wei	哈佛教育学院 (GSE) 学生, Solora Path创始人	1. AI素养: 不应仅针对工程师, 应教育学生/教师/大众。2. 互补模型: 借鉴美国创新、中国规模与加拿大信任。
詹妮弗·图尔柳克 (Jennifer Turliuk)	CIPPIC研究员, MIT MBA/斯隆研究员, JD候选人	1. 无全球条约: 缺乏如《人权公约》的约束力。2. 现有架构有限: 联合国文件不具约束力/可执行性。3. 《欧盟人工智能法案》限制: 仅约束27个成员国。4. 呼吁新机构: 需要有约束力、可执行的国际机构。
阿夫尼·辛哈 (Avni Sinha)	世界隐私论坛 (World Privacy Forum) 高级研究员, 哈佛肯尼迪学院梅森研究员 (Mason Fellow)	1. 追踪全球战略: 研究各国政府如何响应。2. 响应缓慢: 观察显示对技术速度的响应“非常缓慢”。
薛智中 (George Xue)	哈佛肯尼迪学院校友 / 哈佛创新实验室AI教育研究员	1. 迫在眉睫的危险: “AI的危险迫在眉睫”。2. 代理时代治理: 从聊天机器人向代理的转变改变了治理/安全。
肖易菲 (Ella Xiao)	即将入读哈佛肯尼迪学院MPP / Sky9 Capital分析师	1. 全球南方数字主权: 在采用美/中技术时保留议价能力。2. 小国战略: 小国可以“组团”。3. 后发优势: 全球南方的潜力。

颜有毅 (Jason Yen)	多伦多大学罗特曼管理学院 (Rotman) MBA学生 / YC支持的AI项目实习生	1. 企业权力：“拥有太大的权力”。2. 商业与政策：“非常偏向商业侧”，指出两者间的张力。
-----------------	--	--

II. 活动影像总结

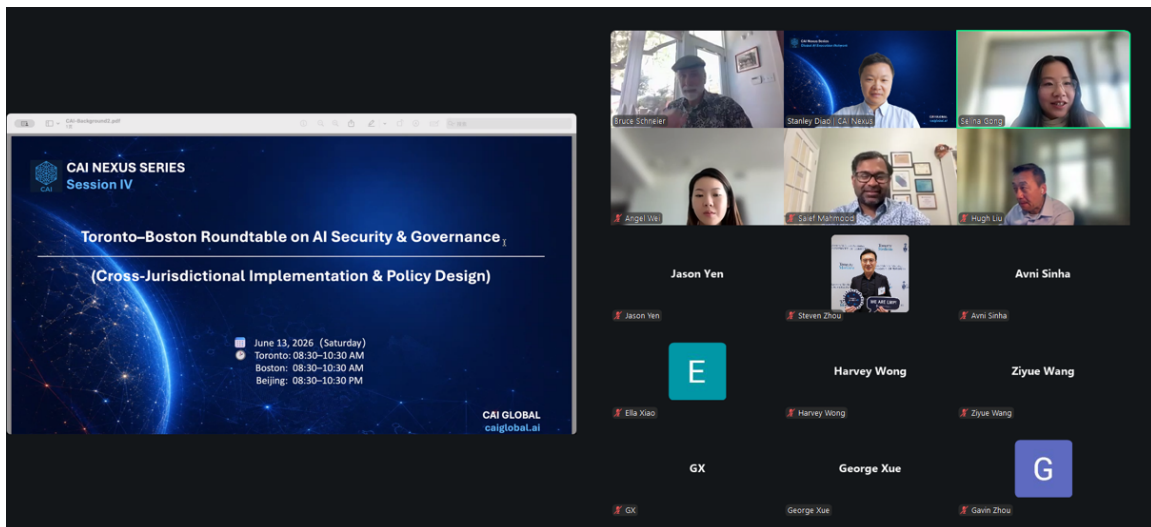
(现场摄影的文字描述)



- 专家分享：布鲁斯·施奈尔 (Bruce Schneier)，哈佛肯尼迪学院 / 伯克曼·克莱因中心研究员。



- 专家分享：赛义夫·马哈茂德 (Saief Mahmood)，加拿大全球事务部高级政策分析师。



- CAI x HFTC系列组织者龚博雅与刁孝力组织AI治理圆桌会议。
- 圆桌会议通过闭门、多节点分布式云连接进行。波士顿和多伦多会场分别由龚博雅与刁孝力主持。连接屏幕生动地捕捉到了波士顿的哈佛学者、渥太华的政策专家、多伦多实验室的核心科学家与亚洲数字经济代表之间的实时碰撞。

III. 参与者所属机构 - 非背书声明

以下机构仅代表参与者的学术或职业背景，不代表其所属机构对本次会议内容或观点的官方立场：

- 加拿大-亚洲人工智能与数字创新理事会 (CAI)
- 哈佛前沿技术俱乐部 (HFTC，暂定名)
- 哈佛肯尼迪学院大中华学会 (Greater China Society, 社区合作伙伴)
- 哈佛肯尼迪学院校友会 (大中华分会，社区合作伙伴)
- 上海交通大学多伦多校友会 (社区合作伙伴)
- 哈佛大学 - 伯克曼·克莱因互联网与社会中心

- 哈佛大学 – 阿什民主治理与创新中心
- 多伦多大学 (医学院 / 罗特曼管理学院)
- 麻省理工学院 (MIT)
- 萨缪尔森-格鲁什科加拿大互联网政策与公共利益诊所 (CIPPIC)
- 加拿大全球事务部 (GAC)
- 世界隐私论坛 (WPF)

IV. 参考文献 (References)

- Allen, D., Hubbard, S., Lim, W., Stanger, A., Wagman, S., and Zalesne, K. (2024). A Roadmap for Governing AI: Technology Governance and Power Sharing Liberalism. Allen Lab for Democracy Renovation, Ash Center for Democratic Governance and Innovation, Harvard Kennedy School.
- CAI Nexus Series IV. (2026). Toronto-Boston Roundtable on AI Security and Governance: Jointly organized by Canada-Asia Council for Artificial Intelligence and Digital Innovation (CAI) x Harvard Frontier Techne Club (HFTC).
- Carvão, P. (2025). The Dual Imperative: Innovation and Regulation in the AI Era. M-RCBG Associate Working Paper Series No. 250, Harvard Advanced Leadership Initiative, Cambridge, USA.
- European Parliament. (2024). Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence (Artificial Intelligence Act). Brussels: European Union.
- Government of Canada. (2024). Pan-Canadian Artificial Intelligence Strategy: AI for All. Ottawa: Innovation, Science and Economic Development Canada.
- Government of China. (2023). Interim Measures for the Management of Generative AI Services. Beijing: Cyberspace Administration of China.
- Harvard Law Review. (2025). Chapter Three: Co-Governance and the Future of AI Regulation. Harvard Law Review, 138(6), 1609 – 1633.

AI Governance White Paper

- Infocomm Media Development Authority (IMDA). (2026). Governance Framework for Agentic AI. Singapore: Government of Singapore.
- United Nations. (2026). Global Dialogue on AI Governance: Summary of Multilateral Consultations. New York: United Nations.
- White House. (2023). Executive Order 14110: Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence. Washington, DC: The White House.