

动荡的 AI 时代已至：我们当下的选择，将决定人类未来

来源：Gates Notes

作者：比尔·盖茨，盖茨基金会主席、突破能源（Breakthrough Energy）创始人

2026 年 8 月 26 日

划重点：

科技的巨轮已然启动，我们需要一份清晰的蓝图，确保利大于弊。

转向 AI 时代的这场大变革，将是人类历史上最动荡不安的时期之一。

此时此刻，我们的准备还远远不够。

但只要走对每一步，AI 就能成为一股向善的力量，让每一个人都过得更好。

我这一生只干过两份工作。第一份是在微软，我致力于通过软件开发赋予人们力量；第二份始于 2008 年，我全职投入慈善，将自己在微软积累的财富回馈社会，希望让世界变得更健康、教育更好、更加公平。这也是我余生将倾尽全力去做的事。

这两段经历，构成了我审视人工智能（AI）的独特视角。13 岁那年我第一次接触电脑，就被“让机器拥有智能、完成当时只有人类才能做的事”这个想法深深吸引。尽管“人工智能”这个词在我出生前后就已出现，但直到近十年，这项技术才迎来爆发式的飞跃。如今它的能力令人惊叹，且正以不可思议的速度迭代更新——人类历史上第一次，AI 不仅能够替代，甚至开始超越人类的认知能力。

AI 将是史上伟大的“平权器”，抑或极恶的“不公源头”？

在公平性方面，AI 要么成为人类有史以来最伟大的“平权器”，要么成为最恶劣的不公源头。我们面临的挑战是史诗级的。即便在最理想的状况下，向 AI 时代的过渡也将是人类历史上最为动荡的时期之一。

我们该如何利用这项技术让世界变得更公平，而不是加剧贫富分化？

我们又该如何保护那些最容易受 AI 伤害的脆弱群体——包括那些失去谋生手段、失去对未来掌控感的人？

回答并落实这些问题，应当是当今世界的重中之重。如果全球能采取正确的举措，AI 将成为向善的力量，福泽每一个人。

遗憾的是，我们现在并没有做好准备。我没有看到各国领导人、专家和社区在积极应对这些挑战。对于如何平稳迈入 AI 时代，全球至今缺乏一份明确的规划。

造成这种现象的部分原因，是许多评论家低估了 AI 即将带来的冲击力。在我看来，主要有以下几个认知误区：

1、AI 目前还会犯错。 就在不久前，AI 连简单的数独都解不开，甚至数不清“strawberry”（草莓）这个词里有几个字母“R”。因此人们很难想象它能取代人类认知。然而，随着研究人员开发出能够自我核查、自我演进的模型，“可靠性”问题正在被快速解决。很快，AI 在许多任务上的表现就会大幅超越人类。

2、拿历史上的技术革新做类比，具有误导性。 我们从未遇到过一种既能被快速普及、又能像人类一样思考和行动的技术。当个人电脑（PC）问世时，它花了整整 20 年才彻底改变我们的工作方式——因为软件需要开发、硬件需要降价，人们还需要学习如何使用工具并将其融入业务流程。但 AI 不同，它直接运行在我们现有的设备上，使用的是自然语言。不是我们要去适应 AI，而是 AI 在适应我们。它能和人类员工看同一段培训视频，直接从现有的数据中汲取知识。

这里我想说明一下我可能存在的局限或偏见。我从科技行业获得了巨大的利益，尽管我的投资组合已经相当多元化，但我与科技界依然存在财务纽带。同时，我也以盖茨基金会主席的身份与微软及其他 AI 公司合作，试图确保 AI 的部署能真正造福全球人民。

不过，我对 AI 的看法绝非出于个人获利动机。我所有投资（包括科技相关投资）产生的利润，都将投入盖茨基金会用于解决全球不平等问题。当然，这是否会影响我的视角，读者自可评判。

这次真的不一样了

在我的记忆中，我总是希望创新能再快一点。但面对 AI，我的内心却更为复杂。

我们需要时间，来为即将到来的社会、政治和经济剧变做准备。

我多么希望世界能够迅速享受 AI 带来的红利，同时尽可能推迟它引发的问题。然而，机遇与危机却同时降临了。最需要时间准备的，往往是那些能力最弱的群体——比如被机器人取代的会计，或是工作被每小时 10 美元的机器人抢走的、原本时薪 20 美元的工人。

许多观察者认为，这次技术转型与以往别无二致。他们举例说，美国当年也经历过劳动力从农业向白领办公室的转移。然而，那次转变历经了几代人，并且创造了需要人类认知能力的新岗位。而这一次，技术直接替代的就是人类的“认知能力”。

由于 AI 具备看、听、说、推理的能力，且未来做体力活也会像人类一样顺畅，它的影响绝不仅限于单一行业。法律、客服、医疗、软件、制造……AI 将全面渗透。这种冲击将在短短十年内快速袭来，而不是几代人的漫长过程。虽然也会产生一些新岗位，但如果没有配套的政策引导，新岗位的数量将远少于当下。

如果有人能提出一份可行的计划来放缓全球 AI 的发展步伐，我大概会支持。但现实是，这不可能发生——地缘政治和经济利益的驱动力太大，全球都在全速推进。

为了最大化这项前所未有技术的正面效应，并最小化其负面影响，我们必须清醒地认识到风险与红利。

转型期面临的三大核心风险

未来我会针对这些风险撰文详述，今天先简要谈谈。

风险一：大量岗位将永远消失

1933 年大萧条时期，美国的失业率约为 25%，并在接下来的近十年里维持在两位数。直到需求、投资和增长复苏，就业才得以恢复。

AI 带来的失业潮可能达不到这个比例，但它的冲击**绝不会随着经济周期结束而自然消失**。受威胁最大的往往是初级和中级岗位，而新创造的岗位大部分都需要耗费多年才能学会的技能。

白领岗位已然遭受波及。在生成式 AI 广泛应用后，极易被取代的岗位中，年轻员工的就业率显著下滑，而年长员工受影响相对较小。

我认为这一趋势还会延续，且绝不局限于少数几个行业。线上和电话客服、销售、软件工程、法律助理等岗位可能首当其冲；但随着 AI 接管评估贷款申请、数据分析乃至急诊分诊等高阶任务，这种颠覆将延伸至更远领域。只要某些环节（如架构设计）依然由人类胜任，软件工程等少数领域会因成本下降而催生新需求，净失业人数会相对较少。

蓝领岗位同样无法幸免。虽然具身机器人（Robotics）的发展略滞后于 AI，但其成本终将大幅下降。我接触过的很多美国人还没意识到灵巧型机器人进化得有多快，因为许多前沿研发发生在海外（主要是中国）。或者他们被网上那些动作笨拙的机器人搞笑视频误导了。在我看来，到这十年结束时，“聪明”的机器人就将在建筑、餐饮住宿等行业与人类展开体力活的竞争。

“机器人+AI”的结合将引发恶性循环：一家企业采用了 AI 并降价，竞争对手就会面临巨大的跟风压力；如果老牌企业不上马，初创公司就会取而代之。许多人被迫转行，但失业、再培训、再找工作的阵痛是巨大的。在市场力量推动下，这种普及只会越来越快。除非我们人工干预，否则优质岗位将越来越少，而技术红利只会向极少数人集中。

我格外担心年轻人，他们步入社会时面对的将是一个初级岗位骤减的职场。作为 AI 最活跃的用户，他们亲眼见证了 AI 的能力和进化速度，因此对未来充满焦虑，这丝毫不令人意外。

最大的转变，将在 AI 能够提供“近乎零差错”的工作成果时到来。届时，它无需人类监督即可自主运转，企业在经济利益驱使下会毫不犹豫地放手让 AI 去做。

这将从根本上颠覆我们对工作、收入和经济安全感（Economic Security）的认知。如果工作的人变少了，或者工作时长变短了，一套建立在“充分就业”基础上的经济体制该如何运转？

在资本主义社会，就业不仅是绝大多数人获取生活必需品的来源，更是尊严和社交纽带的核心。当一个社区出现高失业率时，连锁反应是毁灭性的——研究表明，在美国部分地区，工厂倒闭直接导致了因吸食阿片类药物过量而死亡的人数激增。现在想象一下，如果这种压力同时施加在全国的白领和蓝领身上，会发生什么？

我们必须从现在开始思考如何减少岗位流失，让所有人都能分享 AI 创造的繁荣。等人们已经失业或隐性失业了再行动，就太迟了。AI 是对传统经济组织形式的结构性挑战，需要全新的思考与行动。

风险二：AI 将赋予坏人（甚至 AI 自身）更大的破坏力

早在 AI 普及前，网上就充斥着制作炸弹、生物武器乃至计算机病毒的信息。而 AI 不仅让获取这些信息变得极门槛化，更让将其转化为行动变得轻而易举。即便是技术平平的犯罪分子，也能对个人、企业甚至政府发起大规模精准攻击。

基于 AI 的诈骗、虚假信息、深度伪造（Deepfakes）和无死角监控，将是普通人在日常生活中感受最深的伤害。

AI 也被用于网络攻击。我认识的最顶尖的网络安全专家对未来几年感到担忧——攻击者获取强大新工具的速度，远快于防御方修补漏洞的速度。毕竟，能帮企业找漏洞修补的 AI 模型，同样能帮罪犯利用该漏洞发起攻击。发动一次攻击所需的资源正在大幅降低，而我们目前无法将“恶意利用”与“善意使用”彻底剥离。

想想那些脆弱的基础设施吧：医院、金融机构、供水系统、电网、社会福利发放系统。一旦遭受攻击，受害的将是普通患者、客户和领取救济金的弱势群体。

生物恐怖主义同理。尽管 AI 能加速救命药物和疫苗的研发，但它也能让人更容易设计出致死的新病毒。正向与危险的能力往往是一体两面，这是一个全球性难题。

上述风险谈的是 AI 如何为原本边缘的坏人“赋能”。但与此同时，这些工具也会让权力向既有的权力中心进一步集中。例如，自主武器系统将使政府在无需人类参与决定的情况下使用致命武力；监控和操纵公众舆论将变得更便宜、更高效。

最终，利用 AI 伤害人类的可能不仅是人或机构，还有 AI 系统本身。目前的 AI 系统偶尔就会做出违背设计者初衷的事。这项技术的进化超乎所有人预想，随着模型越来越强大，它们可能会做出违背人类利益的行为，导致人类失去控制权。关于这一点，我未来会专门阐述。

风险三：AI 可能阻碍儿童成长，甚至取代真实的人际关系

我在西雅图长大时，除了几个志同道合的小伙伴，朋友并不多。后来在母亲的大力帮助和自己的努力下，我才逐渐学会社交，懂得如何与不同的人打交道。到了 70 岁的今天，我依然受用于当年学到的社交课。

但我怀疑，如果当年我有一个 AI 伴侣，我还会不会付诸那些努力。

AI 伴侣总是用你最舒服的方式和你说话，永远不会把你逼出舒适区。它们随叫随到，永远不会对你发脾气。这使得它们极具成瘾性，很可能剥夺我们从真实人际交互中学到的宝贵经验。

虽然相关研究仍处于起步阶段，但已经出现了值得警惕的信号。例如，斯坦福大学和卡耐基梅隆大学对 1100 多名 AI 伴侣用户进行的研究发现：社交圈越小的人，越倾向于向聊天机器人寻求陪伴；而这种感情依赖越深、使用越频繁，他们的心理状态反而越糟糕。

这种影响可能是终身的。乔纳森·海特（Jonathan Haidt）在《焦虑的一代》（*The Anxious Generation*）一书中探讨了社交媒体对青少年的危害，而这一点在 AI 身上体现得更为剧烈。他写道：

“就像小树经受风吹才能扎根一样，适度暴露在微小风险中的孩子，长大后能镇定应对更大的风险。相反，在温室里长大的孩子，尚未成年就可能被焦虑压垮。”

一个被设计成“绝不让你生气”的 AI 伴侣，就是一个巨大的温室。

我们才刚刚开始理解互联网（尤其是社交媒体）对青少年发育的危害——强迫性使用、睡眠紊乱、网络霸凌以及不良内容接触。AI 因其更强的说服力和沉浸感，可能会放大这些风险。我们绝不能再等一代人才去重视。包括澳大利亚、英国和挪威在内的国家正在出台未成年人网络保护措施；而中国做得最为彻底——其法规明确限制 AI 伴侣应用，严禁培养情感依赖的设计，并禁止向未成年人提供虚拟亲属和虚拟恋人服务。

除了心理健康，我还担心 AI 对教育的冲击。讽刺的是，这个能让人学到比以往任何时候都多的工具，也可能导致许多人学到的东西变少了。一项初步调查显示，频繁使用 AI 与批判性思维能力的下降呈相关性，且对年轻人影响更为明显。

在深度伪造和定制化虚假信息泛滥的时代，丧失批判性思维是极其危险的。分辨真伪，已成为现代人的必备生存技能。

面对这些心理与社会问题，界限划在哪里尚无定论。在某些情况下，AI 确实能帮人们改善现实中的人际关系；对于孤立无援的老人或行动不便者，AI 伴侣也提供了宝贵的陪伴。但无论界限划在哪里，都应当是我们主动、理智做出的选择。

迎难而上：AI 带来的超级红利

人们常说，我们总是高估一项技术在短期内的变化，而低估它在长期带来的变革。

但在 AI 这件事上，我看到了不一样的现象：有人只看到好的一面，忽视了巨大的风险；有人则走向另一个极端，因恐惧风险而完全无视其巨大的潜力。

我们需要兼具两种心态：既要必须控制的危害保持深刻的警惕，又要对能造福大众的红利保持理性的乐观。

最大化 AI 的红利，与最小化其伤害同等重要。如果大众能直观看到 AI 如何让生活变得轻松，就能建立起应对转型阵痛所需的信任。反之，如果 AI 给大多数人带来的第一印象是“抢

走饭碗”，怀疑者就会彻底摒弃它，届时再想兑现技术的红利将难上加难。这也是为什么政府、医疗等行业以及 AI 公司现在就必须携手合作的原因。

凭借跨学科知识的综合能力，AI 能够加速攻克人类最棘手的技术挑战：提供可靠的清洁能源、应对气候变化、保障粮食安全、消除致命疾病等。在癌症治疗或核能领域的科研人员，可以利用 AI 检索海量文献，发现人类可能忽略的规律，锁定最有希望的实验路线。当“智力”不再是稀缺资源时，小公司也能与拥有巨额研发预算的大机构并驾齐驱，全球的研发与创新将被全面提速。

具体来看，以下领域将率先迎来变革：

医疗健康： 美国许多小型医院缺乏能在紧急关头快速诊断的专家。在这些地方，AI 能及时发现心肌梗死，避免悲剧和高昂的医疗开支。例如软件 Viz.ai 能够通过分析影像检测中风等急症，协助医疗团队协同救治，目前已被全美近 2000 家医院采用。AI 还能协助全科医生精准诊断、跟进患者课后康复，并帮患者解读复杂的检查报告和服药日程。

农业（低收入国家见效最快）： 在大多数低收入国家，农户得不到准确的气候预报，不知道该选什么种子、如何防治病虫害或改善土壤。在人口增长和气候变化的双重压力下，他们比以往任何时候都更需要帮助。借由 AI，贫困地区的农民很快就能获得比当今最富裕农场主还要精准的农业建议，大幅提升粮食产量。

政府公共服务： 在美国，许多家庭在申请医疗保险、助学金或食品补助时，常常面对繁复的官僚表格而望而却步。AI 可以大幅简化这一流程，让弱势群体更快获得救助，提升政府运转效率。

心理健康辅助： 绝大多数社区都面临心理咨询师和戒瘾专家匮乏的问题。在严格保护隐私的前提下，AI 能识别早期预警信号，提供循证应对策略，并与人类专家配合提供更高效的治疗。

教育赋能： AI 能将教师从繁重事务中解放出来，投入更多精力进行一对一或辅导小班教学。对学生而言，好的 AI 工具会保留所谓的“有益挫折”（Productive Struggle）——即构建深度理解所需的认知思考过程。当学生首次接触新概念时，AI 提供详尽解答；而在检测理解程度时，它会扣留答案，引导学生独立思考。

总而言之，这些突破将使日常生活变得更轻松、更实惠，且不再受限于个人的收入或社会关系。

AI 能让普通人和小微企业拥有过去只有高薪专家或庞大团队才能提供的能力，同时让产品和服务更优质、更便宜。它能帮助残障人士更独立地生活，让有创意的劳动者和创业者实现超越以往的成就。

最重要的是，它能把人们从填表、官僚流程、筛选信息等琐事中解脱出来。当这些微小的改变乘以数以亿计的人口，将汇聚成巨大的社会进步：更多人在需要时获得良言，拥有更大的自由去追求理想的生活。

盖茨基金会的行动

AI “能够”改善所有人的生活，但这不会自动发生。我们必须有意识地引导，确保它惠及每个人，而非仅仅服务于少数富人。这就需要政府和慈善机构发挥强有力的作用。

盖茨基金会在其 20 年存续期内还剩 19 年，期间将花完剩余的 2000 亿美元。AI 将通过加速艾滋病、结核病、疟疾和营养不良的疫苗与药物研发，助力我们实现宏伟目标。基金会的目标包括：将全球每年死亡的儿童人数在 2000 至 2024 年的基础上再减半。我们在健康、农业和教育方面的所有工作，都将全面拥抱 AI。

下个月，我将在基金会的年度《目标守护者》（Goalkeepers）报告中阐述更多细节——包括确保 AI 模型支持中低收入国家的本土语言，而不局限于富裕国家的语种。包括 OpenAI、Anthropic、Google 和微软在内的顶级 AI 公司正在与我们展开合作。

破解难题：世界需要一份“应对蓝图”

部分 AI 公司能针对自身技术带来的挑战提出解决方案，这很好，但我们不能指望它们来领头。有些问题超出了它们的专业范畴；况且在民主社会中，由企业决定人类的未来并不合适。

解决方案必须通过公开的民主程序来制定，涵盖民选官员、决策者、教育工作者、医护人员和社区领袖。面对数以百万计生活受冲击的人们，我们需要一个更强大、更具弹性的社会安全网。

这些方案应当建立在我们对 AI 根本性问题的回答之上——包括当机器比我们更聪明时，我们该如何守住“人性”？终身在思考“何为人”的宗教领袖们在此能发挥关键作用。教皇良十四世（Pope Leo XIV）关于 AI 的通谕《在人工智能时代守护人类》（*On Safeguarding the Human Person in the Time of Artificial Intelligence*）令我印象深刻，它为未来的工作奠定了坚实基础。

围绕如何确保 AI 利大于弊，我提出以下三点先行倡议：

1. 建立全新的转型管理体制（重中之重）

当务之急，是构建一套应对 AI 的国内与国际治理框架。

现有的机构没有一个是为处理传播如此之快、影响如此之广的技术而设计的，我们需要建立全新的机构。其工程之浩大，怎么强调都不为过。911 事件后，美国政府为了国土安全进行了自二战以来最大规模的部门重组。而 AI 所需的工作量远超于此——它同时影响着国家安全、就业、教育、税收、能源、选举、公共卫生、金融、执法和交通等各个领域。

这些领域交织在一起，传统的官僚体系无法应对。劳工部门懂就业，但不懂安全风险；市场监管部门懂垄断，但不懂 AI 对青少年的心理影响。如果不打破壁垒，各机构只会盲人摸象，而 AI 的影响却是牵一发而动全身的。

在国家层面，各国需要能够统筹各部门优先事项的机构，确保无死角覆盖风险。而在跨国层面，由于风险无国界，必须同步建立国际组织。

这个全新的国际组织可以借鉴现有的模式：核武器检查机制、国际民航法规、保护臭氧层的国际协定。新的全球 AI 组织需要兼具这些要素。

考虑到政府办事效率低下，以及国家内部和国家之间的政治对立，人们有理由怀疑现有的国际机构能否胜任。但这需要美中两国开展一定程度的合作。

我们没有慢慢来的资本。国家领导人应定期召集经济学家、技术专家、劳工专家和企业界人士，明确现有机制的失灵之处，探讨需要赋予什么样的新职权。掌握核心供应链和顶尖开发者的国家，应当在市场竞争压力加剧之前，尽快坐下来制定共同规范。

搭建这个框架需要数年时间，正因如此，我们必须**立刻行动**。

2. 设立“人类保留区” (Human Reserved) 岗位

2020 年，我的父亲因阿尔茨海默病去世。在他生命的最后阶段，日夜照料他的是一群充满爱心的护工。即便父亲表达困难，他们也能敏锐地感知他的需求。那种关怀中包含着不可替代的人性温度——没有任何机器人能够、也不应该取代这份工作。

每当讨论哪些岗位会消失、哪些会保留时，我就会想起那群护工。我认为，随着 AI 和机器人的演进，我们应当划定某些领域，**指定仅由人类来完成**。我称之为“人类保留区” (Human Reserved) 。

我喜欢这个词，因为它让我想起自然保护区——那些地方本可以修路盖楼，但为了保护珍贵的生态，我们选择不破坏。

划定“人类保留区”可能是出于经济考量：避免机器全面接管导致大量难以转行的人失业（你不能指望一个干了一辈子建筑的 55 岁工人转行去养老院并从中获得成就感）。有时则是出于伦理和情感：比如在医疗领域，技术上完全可以让机器人告知患者患上了绝症，但我们绝不应该这么做。

“人类保留区”的边界将动态演变。例如，我们可以选择在某些领域保留人工，通过承诺保留部分岗位，在几年或几十年内循序渐进地引入 AI。而在教育和心理健康等领域，则可以采取“人类主导 + AI 辅助”的复合模式。

边界因地制宜。有些国家可能坚持养老必须由人类完成；而像日本这样面临劳动力萎缩、缺少年轻人照顾老人的国家，可能非常欢迎护理机器人的加入。

这个概念也引发了一连串需要公众共同探讨的问题：*谁来决定保留什么？标准是什么？如何防止企业违规私用机器人？当一国允许机器人生产而另一国禁止时，国际贸易该如何处理？*

3. 重塑资本与劳动力的税收平衡

当劳动者被迫转向新岗位时，他们需要再培训和社交安全网的支持。但如果工作时间减少、缴纳的个人所得税降低，政府就会在最需要资金的时候面临财政收入锐减。

我主张对 AI Token（算力使用单位）和机器人征税。

在现行税制下，雇佣员工需要缴纳雇主薪酬税（Payroll Taxes）；但如果购买一台机器人，通常可以作为业务费用直接抵税。这种税制实际上是在变相鼓励企业用机器取代人工。

征税不仅能适度放缓企业逃离人工劳动的脚步，还能再培训和社会安全网筹集资金。税收应当进行精准定向，切勿阻碍 AI 在降低医药和教育成本方面的正面应用。

批评者认为这在纯经济学意义上效率不高，但他们忽视了工作对个人和社会的广义价值。在技术革新带来巨大生产力提升的背景下，为了保障大众就业，我们完全承受得起这一点效率上的让步。

几年前我就提出了“机器人税”，当时大多数人觉得这是个怪异的想法，但我至今仍是坚定的倡导者。它虽非万灵药，却是明智应对方案中不可或缺的一环。

无论通过何种方式筹集资金，援助都必须精准触达最需要的人——包括因 AI 失业的工人、工时与工资下降的群体，以及受冲击最严重的社区。

我正在做的事

我将倾尽我的声音和时间，将“AI 与社会公平”推向公共议题的中心。每次去华盛顿或会见全球领导人时，我都会阐明这一立场；在与 AI 模型开发者交流时，这也会是核心议题。

我将全力推动上述国家与国际治理框架的建立。盖茨基金会将致力于推动 AI 在贫困地区（包括非洲）的善意应用；突破能源公司（Breakthrough Energy）将利用 AI 开发廉价清洁能源，助力解决气候危机。此外，我也会定期撰文分享关于 AI 的思考。

我向各国领导人呼吁：

在失业率飙升、社区受到伤害、公众信任崩塌之前，你们现在就有机会采取行动。

你们可以确保政府跨部门统筹治理，而不是将其切割成官僚割据的领地。

你们可以确保 AI 惠及每一个人，并与其他国家携手应对这一全球性挑战。

最后，我将努力扩大参与这场大讨论的圈子。这不应仅仅是少数技术专家的游戏，更应涵盖工人、即将在 AI 时代步入职场的大学生、社区领袖、宗教团体、家长和教育工作者——他们的声音往往被忽视，但他们对这场技术转型将如何影响普通人的生活，拥有最深刻的洞察。

我们该如何确保 AI 的红利触达那些尚未拥有财富和影响力的人？

当岗位被颠覆时，我们该如何巩固社会安全网，帮助劳动者与社区繁荣下去？

公共机构应当如何调适以适应新时代？

最重要的是，在这场前所未有的变革中，我们该如何守住我们的“人性”？

这项前所未有的技术，需要人类做出前所未有的全球响应。如果我们做对了，人类将获得无法估量的巨额回报，世界也将变得更加公平。

我几乎无时无刻不在思考 AI——并非因为我掌握了所有答案，而是因为它抛出的问题太沉重，绝不能仅仅交给一小撮技术专家来决定。学术界、企业界、政府和民间社会的领袖们，都必须在塑造未来的过程中担起责任。

来源: Gates Notes

作者: 比尔·盖茨

翻译整理: 创业邦

原文链接: <https://www.gatesnotes.com/a-turbulent-ai-era-and-critical-choices-to-make>