

生成式人工智能革命八阶段： 现在你在哪儿？

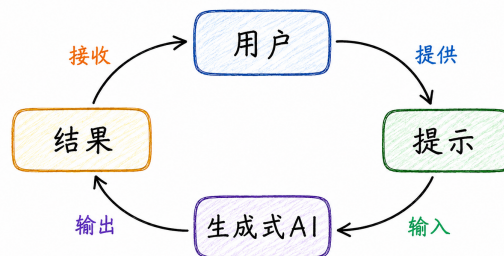
王士弘 (Paul S. Wang)

2026 年 8 月 20 日

要讲清楚生成式人工智能（GAI）究竟是什么，并不容易。简单的说，传统 AI 是“专才”，GAI 是“通才”。

传统 AI 系统为完成某一项特定的智能任务而构建，例如人脸识别、语音辨识、指纹匹配，或是自动驾驶。而 GAI 系统则截然不同。它们被设计为通过自然语言与人类对话，不仅能理解你的问题，更能代你生成内容、协调多个任务、调用传统 AI 工具，像一个真正的智能助手那样为你“办事”。

你与它交谈的每一个提问或评论，被称为提示词 (*prompt*)。用提示词来说明你提出的问题或任务，而它就会据此进行回答与操作。



GAI 提示循环

GAI 系统基于大规模语言模型（LLM）构建，它采用一种叫做 Transformer 的内部算法，利用提示词作为输入来生成输出。提示词既可以通过键盘输入（更精确），也可以通过语音输入（更便捷）。这些正是其技术突破的特性。

这里有必要明确一点：不要将 GAI 与 AGI（通用人工智能）混为一谈——后者就像建立火星殖民地一样，目前还只是一个遥不可及的梦想。相比之下，GAI 是一种智能计算系统，其运作本质上就是不断执行提示循环。

在对话过程中，GAI 会像与人面对面自然交谈一样，让语意始终保持在上下文中。我们都知道一遍又一遍地重复自己说过的话有多么烦人。

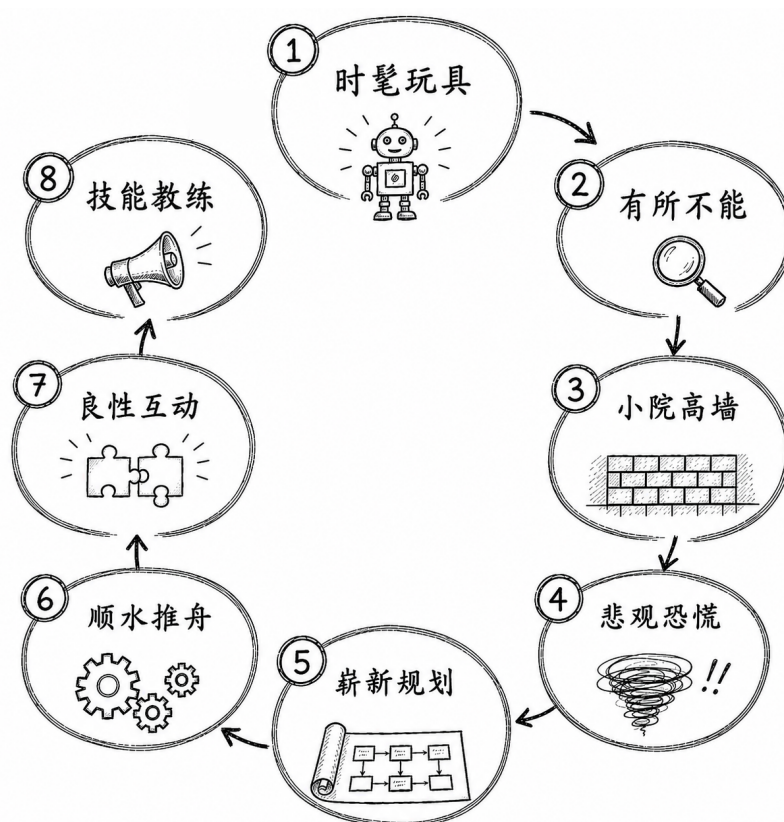
更关键的是，GAI 不仅“知道”人类已有的知识，它还能创作和生成全新内容：从一首诗、一段代码，到一份商业计划书、一幅设计草图。这种知晓与创造能力的双重叠加，使其威力呈指数级放大。以至于今天，从教育、医疗到制造、金融，从专业工作到日常生活，

几乎每一个角落都在被它赋能和重塑。现代搜索引擎能够识别查询中的提示词，并自动为您调用 GAI。换句话说，**搜索即对话、对话即服务**的范式已经落地。

GAI 革命中有技术飞跃，也有人们心理适应的过山车。如果仅仅把 GAI 看作一次技术升级，我们就忽略了它更深层的影响：**对人类自身独特性信心的挑战**。当一台机器突然可以在几秒钟内与你对话、写诗、推理、编程、设计、分析数据，甚至在某些维度上超越多数专业人士时，人类的反应绝不仅仅是换一种工具那么简单。我们会焦虑，担心自己被取代；会兴奋，惊叹于前所未有的能力延伸；也会困惑，不知道该如何适应，亦不把握自己该扮演的角色。

大多数人的 GAI 之旅，并非一夜之间从初学到专家。而是会经历一系列心理与学习的阶段。从一开始的震惊与抗拒，到试探，再到有效利用，最终走向深度融合。这一过程，既是对新技术的适应，也是对自我的重新认识。

因此本文提出“GAI 八阶段模型”，为读者展示从“莫名其妙”到“传道授业”的心理过程。此模型结合每个阶段的行为与心态来帮你了解：此刻的你，走到了哪里？



GAI 适应八阶段模型

在阅读接下来的八个适应阶段时，请在每个阶段停留片刻，回答我们给出的扪心自问。不仅问自己，也可以问你的朋友、同事和团队成员。因为这场革命最有趣的地方在于：每个人会取得自己的进展，而了解彼此的位置，正是共同前行的起步。

第一阶段：莫名其妙——“时髦玩具”

内心感受：不明其中奥妙，把 GAI 当作一种在线的新奇玩意儿或聚会上的小把戏。

行为典型：你把 GAI 当成一只会说话的开心小狗。有趣，但仅此而已。你让它“用老鼠的口吻写一首关于猫的诗”，或者“生成一张穿着西装的土豆的搞笑图片”。它逗你开怀大笑，你也乐于在朋友圈或微信群里分享这些 AI 神回复的截图。但它跟正经事毫无关系。

在这个阶段，你可能已经注册了文心一言、通义千问或 ChatGPT 的免费账号，甚至也尝试过 Midjourney 或国内的通义万相生成过几张图。但你用它们的方式，基本停留在看看它有多聪明的层面。你的每一次提问，都更像是在逗弄一个新鲜物种，而不是在解决真实问题。

我们都会从玩开始，这并非幼稚，而是人性使然。面对一项颠覆性技术，人类的自然反应不是我该怎么用它改变世界，而是这玩意儿到底是干啥用的？这种好奇和把玩，恰恰是我们建立直观认知的最重要途径。得先摸清楚这头大象的鼻子有多灵活，才会考虑如何让它干活儿。

扪心自问：“你有没有试过让 GAI 写一首关于你老板的、用东北话押韵的说唱歌曲？或者让它编一个孙悟空大战吴三桂的荒诞故事？”

如果试过，恭喜你，你正处于起点。很多人玩几次 GAI，新鲜感一过，就把它忘在脑后，这样就错失了接下来那扇即将打开的门。

在中国，玩具阶段的场景已经比两年前丰富得多。有人会让 Kimi 或豆包写一首“赛博朋克风格的唐诗”，用讯飞星火生成一个“如果诸葛亮穿越来做产品经理”的职场段子，或者让智谱清言用四川话给你讲一个睡前故事。这是 GAI 生态中的吃瓜群众，他们中的绝大多数，不会想要迈出下一步。

第一到第二阶段的跨越，往往来自一个偶然的瞬间，你突然发现，这个玩具竟然真的帮你解决了一个小麻烦。那一刻，玩具变成了工具。

第二阶段：测试局限——“有所不能”

内心感受：故意给 GAI 出难题或刁钻的专业问题，就为看它翻车，然后长舒一口气。

行为典型：你内心的怀疑终于浮现了。“那些花里胡哨的玩意儿不过是包装炒作，我倒要看看你是不是真有两下子。”你带着这种心态，翻出你工作中最刁钻、最冷门、最让新同事崩溃的那个专业难题，一个只有资深工程师才懂的边缘案例，或者一个涉及最新政策法规的复杂情境。你把它抛给 GAI，带着一种考官出题的优越感。

然后它果然错了。要么是计算出了偏差，要么是引用了过时的条款，要么干脆幻觉发作，自信满满地编造了一个压根不存在的学术名词或案例。你盯着屏幕，嘴角上扬，内心升起一种奇异的满足感：“哈！我就说吧！它根本不懂我这一行。”这一刻，你如释重负。你心里那个被取代的隐忧，暂时被安抚了。你把截图发到同事群里，配文“AI 也不过如此”，收获一波认同的点赞。

天然的自保让我们热衷于“抓包。”面对一项可能威胁到我们专业身份的技术，我们本能地选择寻找它的弱点来证明自己的不可替代性。每一次抓住 GAI 的失误，都是在对自己说：“看，我十几年的经验积累，哪是一个模型能轻易复制的？”这种心理防御在中国职场中尤为普遍。从程序员到律师，从设计师到会计师，几乎每个行业都经历过这场“集体松了一口气”的时刻。

但这里藏着一个危险的认知偏差：你测试的往往是 GAI 当前的极限，而忽视了它迭代的速度。你拿现在某个模型的失误来证明“它永远不行”，却忘记了过去 18 个月它能力跃升的曲线有多陡峭。今天的缺点，三个月后可能就不见了。

扪心自问：“当你故意刁难 GAI，而它果然犯错的时候，你是否有一种饭碗稳了的安全感？”

如果是，你正处在第二阶段，也是大多懂行人士容易长期卡住的阶段。这个阶段的人往往拥有较高的专业素养和批判性思维，但他们的批判性只对准了 GAI 的短板，却没有同样严格地审视自己的可替代部分。

中国的职场中，第二阶段的抓包游戏有几种常见形态：法律从业者让 GAI 分析一个涉及《数据安全法》和《个人信息保护法》交叉的复杂案例，然后指出它对某条司法解释的理解偏差；程序员扔给它一段罕见的冷门代码，看它能否正确解读和优化；会计师问它关于最新税收优惠政策在某一特殊行业的具体适用条件，然后指出它遗漏了某个地方性的补充通知。每一种抓包都能带来短暂的满足感，但它同时也暴露了一个有趣的事实：只有当你真正深入了解一个领域时，你才能识别出 GAI 的错误。换句话说，能抓包本身就是专业能力的证明。而真正的问题在于这份专业能力是用来拒绝还是引导这项技术。从这里到下一个阶段的跨越，往往发生在一个你没有刻意刁难，却发现 GAI 意外帮了大忙的时刻。那个时刻会迫使你重新审视一个问题：“如果它已经能在我不擅长的领域给我惊喜，那我专门测试它擅不擅长我的领域还有多大意义？”

第三阶段：负隅顽抗——“小院高墙”

内心感受：固执己见、推动职场禁用、或痛斥技术扼杀人类的创造力。

行为典型：此前那种“抓到你了”的短暂快感，正在迅速消散。因为，不管你愿不愿意承认，那个模型确实在变好。你上个月用来刁难它的那道刁钻问题，这个月它竟然答对了七八成。你再想找出一个能“完美翻车”的案例，变得越来越费劲。而正是这种进步本身，让你感到真正的不适。于是现在要全力抵抗，变成挡在那里的一堵墙。

表现形式各有不同：你刻意回避所有 GAI 工具，即使同事提到某个好用功能，你也假装没听见；你在朋友圈、知乎或即刻上发帖，痛陈“AI 正在杀死真正的创造力”、“技术让人类变懒”、“以后没人会真正思考了”，收获一批同频者的共鸣；你甚至开始在团队会议或公司内部推动禁用 GAI 的政策：“为了信息安全”、“为了保证原创性”、“为了防止我们团队变得平庸”是你的主要论据；甚至你会半开玩笑半认真地说：“用 AI 写东西的人，我一般不把他当同行。”为什么我们会“筑墙”？核心不是对技术的理性批判，而是有自我认定的危机。我是一个有能力的人。但 GAI 会让我丢掉工作，那我的价值何在？

这种恐惧驱动下的抵抗，常常披着“守护人类价值”的外衣。但我们需要诚实地区分两种批评：一种是有建设性的边界探讨（比如“哪些创作领域应该保留人的主体性”），另一种是拒绝面对现实的防御性姿态（比如“凡是 AI 做的都是垃圾”）。前者是对话，后者是筑墙。

扪心自问：“你是否在主动回避 GAI 工具，甚至暗自希望你所在的组织直接一刀切地禁止它们？”

如果是，那么你正处在第三阶段。这也很自然，因为是你的职业自尊心在发出警报。但是你打算在这堵墙后面待多久？

“那堵墙”现象在中国不同群体中有不同的表现形式。高校教师群体中，部分教授公开呼吁禁止学生用 AI 写作业，但私下里自己用 AI 写课题申报书的现象并不罕见。这种双重标准，恰恰是身份抵抗的扭曲表现；创意行业中，设计师、文案、插画师等群体主张“AI 作品不应获得版权保护”，其实是对自身创作价值被稀释的焦虑；另外，部分人用“国产模型不如国外”作理由抵抗：“国产的我不信任，国外的用不了，所以干脆不用”。这看似有理但已不符合实情：截至 2026 年中，国产 GAI 在中文理解、政务合规、财务处理等本土化场景上已有优势。

第三到第四阶段的跨越，往往不是因为你想通了或被说服了，而是因为现实压力。你发现竞争对手在用 GAI 比你先交付方案，你发现新入职的年轻同事效率比你高出一截，你发现客户开始用 AI 生成的初稿来要求你修改而不是从头做。当抵抗的成本变得高于接纳的成本时，墙才开始有裂缝。但有趣的是，跨过这道坎的人往往会成为最有辨识力的 GAI 使用者。因为那些真正抵抗过、真正认真思考过什么值得保留的人，比那些从不抵抗就全盘接受的人，更能明智地决定“什么交给 AI，什么留给自己”。

第四阶段：生存危机——“悲观恐慌”

内心感受：心里一沉，一个新版 GAI 模型，五秒钟就能完成你苦练多年的本事。

行为典型：这一天终究来了。

你可能是在某个普通的工作日下午，随手刷到一条新闻：“某国产大模型发布新版本，多项能力达到国际领先水平”。你没有太在意，点开链接随便试了试。随手丢给它一份你花了整整一个周末才搞定的外文合同翻译，或者一段你调试了两天的复杂代码，或者一个你过去要调研三天才能写出来的市场分析框架。

五秒钟后，它给出了答案。

不是那种敷衍的、需要你大改特改的垃圾初稿。而是好到让你后背发凉的那种：结构清晰、逻辑自洽、细节准确，甚至比你做的版本还多覆盖了几个你没考虑到的角度。

这一刻你心凉了。那种感觉不是它好厉害的好奇，不是它犯错了的得意，不是我不喜欢它的抗拒，而是“完了”。

盯着屏幕你脑子里翻涌着几个挥之不去的念头：我花了六年读的学位，到底算什么？我这些年熬夜积累的经验，它几个月就学会了？明年这时候，我这个岗位还在吗？我接下来该往哪儿走？

关掉网页、合上电脑，试图让自己冷静下来。但那个问题像一颗种子，已经扎了根：“如果机器已经能做到这一步，那我的价值究竟是什么？”

恐慌如此剧烈，因为它击中的是我们最深的职业身份根基。这种恐慌有一个特别的触发机制：迭代速度的非线性。很难预测下一个版本又会在哪个领域超车。在某个领域深耕多年专精一门的中坚力量，感受到的冲击最为猛烈。

扪心自问：“最近 GAI 的更新，是否让你暗自害怕自己的职业或学位正在过时？”

如果是，你正在经历第四阶段。这恐慌是清醒不是软弱。那些还不恐慌的人，要么是没亲身体验，要么是未受影响。而你已经看清了。

在中国，恐慌往往比西方更加多元。翻译和外语专业：2026 年国产模型的中英互译质量已接近甚至在某些领域（如政务公文、法律条款）超过中级译员，初级程序员：通义灵码、CodeGeeX 等 GAI 编程助手的普及，让“增删改查”类的基础开发岗位需求显著收缩。恐慌在当今的中国已从个人的心理状态成为公开社交的话题。豆瓣、小红书、知乎上涌现了大量“GAI 焦虑互助帖”，人们分享自己的经历，寻找下一步该学什么的答案。

从四到第五阶段的跨越，是整套模型中最微妙、也是最关键的一步。它的起点不是一次技术更新，而是从“我被取代了”到“我被解放了”的观念转换。这个转换需要你主动把目光从 GAI 能做的转移到不能做的，但不带第二阶段的嘲讽。于是开启一种新的探索：在那些事情里藏着我的新位置？当你用会让我成为什么样的人来评价它，恐慌自然就结束了。

第五阶段：战略转向——“嶄新规划”

内心感受：找到我们能做而 GAI 不能做的是真正的出路。

行为典型：恐慌像一场高烧，它来势汹汹，但也终会退去。而退去之后留下的，不是回到原来的状态，而是一种清醒的冷静。

你不再瘫在椅子上盯着屏幕发呆。你站起来，走到桌前，拿出一支笔、一张白纸，或者打开一个空白文档。你开始写。

你画了一条线。线的左边，列出所有机器可以自动化完成的事情：数据整理与初步分析、标准文档的初稿撰写、信息检索与摘要、代码的框架生成与调试辅助、翻译与基础校对、报表生成与可视化。

线的右边，列出那些模型至今无法真正触碰的领域：人与人之间深层的情感连接与共情；复杂情境下的最终决策责任，“我来签字，我来承担后果”；基于长期信任关系建立的客户和团队纽带；对模糊、矛盾、不完整信息的直觉性判断；将一个领域的概念创造性地用到另一个领域（跨界类比）；对“什么是好的”和“什么是对的”的价值判断，不仅仅是逻辑上的对错，而是伦理和美学的取舍。

你看着右边那一列，第一次意识到：这些不是 GAI 的缺陷，而是人类的领域。过去你从未把它们当作核心竞争力，因为从来没有人能像 GAI 一样在你擅长的左边一栏里跑得这么快。于是你决心改变战略：左边交给工具，右边留给自己。主动设计交给 GAI 的任务，并加强自己的能力。这就是你存在的理由。

为什么画线如此重要？第五阶段的本质是认知重构。从受害者的消极叙事，转向主动作为的积极态度。

这一步需要勇敢的诚实：既不能低估 GAI 的能力（那是第二阶段的自欺），也不能夸大它的威胁（那是第四阶段的恐慌）。你需要像划国界那样，冷静地标出每一块领土的边界，然后决定你的旗帜插在哪里。有趣的是，当你真正完成这规划图之后，你会发现一个意外的事实：左边那栏越大，右边那栏的价值就越高。当所有人都能借助 GAI 高效完成左边的工作时，右边那些不可自动化的人类特质，反而变得比任何时候都更稀缺、更昂贵。

扪心自问：“你是否已经梳理清楚，自己工作中哪些部分可以交给 GAI，从而让自己能全身心投入到高价值的人类技能上去？”

如果是，你已经进入了第五阶段。这个阶段的人不会再去争论“GAI 有没有我厉害”，因为这个问题已经不再重要。他们关心的是另一个问题：“有了 GAI 之后，我能成为什么能力更强的人？”现在的中国职场中，第五阶段的规划呈现出几个本土特色。“新质生产力”的政策框架为第五阶段的思考提供了一个宏观语境：不是强调“技术替代人”，而是“技术赋能人，让人去创造新的价值形态”，这与第五阶段的战略转向一致。许多企业开始组织工作坊训练 GAI 与人的新分工。一些领先的中国企业开始进行岗位职责重写。将传统岗位拆解为可自动化和需人类单元，重新设计岗位说明书和绩效考核体系。

第五到六阶段的标志性事件是你真正地把 GAI 集成到你的工作流中，感受到“1+1 大于 2”的协同效应：不是“它替我做了我的工作”，而是“它让我做了我原来做不到的工作”。那一刻，工具成了你工作的一部分。

第六阶段：有效利用——“顺水推舟”

内心感受：不再聊天对话，而把 GAI 当作一台有严格输入与使用规则的数据处理器。

行为典型：你发现聊天的效率太低。上阶段你画的那张规划图上，左边那一栏（可交给机器的任务）现在到了任务执行阶段。而执行的第一条铁律是：提示词决定结果。

之前：“帮我写一份关于新能源汽车市场分析的 PPT 大纲，大概十页左右。”

现在：

【输入数据】 上传文件为今年 1 至 6 月乘联会销量数据、五家头部企业财报摘要、以及三份行业研报核心观点。

【约束条件】 每页不超过 3 个要点；每个要点不超过 15 字；必须包含竞品对比矩阵；避免任何预测性表述；输出格式为 Markdown 层级列表。

【角色】 一个资深战略咨询顾问，用专业、数据驱动语气。

与聊天不同，你给它明确的任务、清晰的边界、输出的要求。你把提示词格式化，变成一套有结构的指令。你也会开始建立自己的提示词库：针对不同的任务（摘要、分析、生成、转换、审查）有不同的模板，以供大家使用。这是业余与专业者最明显的不同。

你会发觉：GAI 的幻觉和错误不再让你沮丧，而是让你改进指令的信号。每次它出了偏差，你会问自己：“我的提示哪里不够明确？我的约束条件哪里有空隙？我的输出要求哪里模糊了？”毕竟，掌握新工具需要具备新技能。

扪心自问：“你是否正在创造结构化的提示词而不再只是提一些随意的、闲聊式的问题？”

如果你已进入第六阶段，了解 GAI 不是魔法，而是利器。事实上，提示工程已成为一种职业技能。但你也不要过于控制，因为那会扼杀 GAI 的发挥和在互动中涌现的能力。

从六到七阶段的跨越，往往发生在一个放松控制的时刻，你给它一个开放性的挑战，它给你的回应却远远超出你的预期。那一刻，你意识到：最好的方式，不是把机器锁死在轨道上，而是要让它还有可以自由发挥的空间。你开始从工程师变成园丁，不仅建造有效系统还要培育可能性。

第七阶段：人机共舞——“良性互动”

内心感受： 将 GAI 融入日常，让它像打电话或用互联网一样自然。

行为典型： 你的浏览器里有一个永远打开的 GAI 标签页。它已经在那里很久了，习惯了。GAI 不再是个让你兴奋、恐慌、或努力改进提示词的东西。它变成了空气。

你早上打开邮箱，不再从头开始回复每一封邮件。你快速浏览，在脑子里标注要点，然后让 GAI 生成初稿，你花三十秒修改语气和关键措辞，点击发送。整个过程行云流水，以至于你很难说清楚“哪些是我写的，哪些是 AI 写的”，界限已经模糊了。

你在开会时谈到一个你不熟的问题，你不再说“我回去查一下”。你打开对话框，把问题丢进去，很快拿到了一个答案。因为“不知道就问 AI”已经变成了你的本能反应，

在写一份方案时，不再坐在空白文档前苦思。你对着 GAI 口述你的初步想法。可能是一团乱麻，可能是零散的念头，甚至是一些彼此矛盾的思路。然后你看着它整理出来的结构，删掉一些、补充一些、推翻一些、重塑一些。写方案不再是孤独跋涉，而是一场有的放矢的共舞。

你开始在 GAI 的产出中看到自己的影子：因为它已经把你的思维习惯、表达偏好、质量标准带进了它的回应中。那些回答不再陌生，而感觉像你做得更快。

为什么无形是最高的使用状态？第七阶段的本质，是你把 GAI 从新工具到变成身体的延伸。

回想一下，你现在打字的时候，会意识到键盘的存在吗？你上网的时候，会感叹互联网好厉害吗？不会。它们已经退到了意识的背景里而不是你注意的对象。第七阶段就是这个状态。GAI 不再是你使用的东西，而成为你思考的方式。当一个问题出现时，“我可以 AI 帮我处理这部分”已经很自然，就像开灯和关灯一样

在这个阶段，你不再觉得 GAI 威胁你，也不觉得你在驾驭它。你只是和它在一起，各自做各自最擅长的事。

扪心自问： “你的 GAI 标签页是不是整天开着为了那些重活累活，而你就负责最终的判断和拍板？”

如果是。你已成为真正将 GAI 融入日常的人。但要警惕第七阶段特有的陷阱：舒适让你变得懒惰而可能会把一些本应亲自动手、保持手感的事情也交给了它。对此保持警觉，是让你迈向最后一个阶段的阶梯。但第七阶段有一种新的焦虑：“如果我的同事看不出来哪些是我写的、哪些是 AI 写的，那他们看重的是我还是 GAI？”因此，有些人开始主动标注：“这部分是我独立思考的”“这部分是 AI 建议但我做了重要修正”。这种透明化是一种新的职业诚实，也逐渐成为专业社群中的一种不成文的伦理标准。

七到八阶段的跨越，是整套模型中最微妙的一步。它不来自任何技术更新或技能提升，而是来自一个视角的转换：当你不再问 GAI 能为我做什么，而开始问什么是一方单独做不到的，那时你就到达了最后一站。

第八阶段：传道授业——“技能教练”

内心感受： 制定企业准则、为同事创建提示词模板、把其他人一起拉进未来。

行为典型： 你抬起头来。在经历了前七个阶段的漫长旅程之后，你的目光从自己和 GAI 的互动中移开，转向了周围的人群。你看到的是什么？

你看到大多数同事还在第二阶段和第三阶段之间徘徊。他们还在玩抓包游戏，还在抱怨 AI 会毁了这个行业，还在等待某个人来告诉他们该怎么做。

于是你动起来了。你开始整理自己的提示词库，把它从个人笔记变成一份可共享的团队资源。你给每个模板加上注释：什么时候用、什么时候不用、可能出现什么问题、怎么调试。

你在部门会议上主动站出来：“我最近在用 GAI 做某工作，效率提高了很多，我可以给大家分享经验。”

你开始起草一份团队层面的 GAI 使用指南，不是那种禁止使用的政策文件，而是如何安全、有效、负责任地使用的操作手册。你写清楚哪些类型的数据可以喂给模型、哪些绝对不能、输出需要经过怎样的审核流程。

你在公司的内部 wiki 或飞书文档里建了一个 GAI 实用案例库，鼓励大家分享自己的用法。一开始没人响应，你就每周亲自更新两三个案例，渐渐地，别人开始提交自己的案例了。

你甚至开始在这个话题上得到了一些影响力。年轻同事私下叫你 AI 教练，HR 部门邀请你为新员工做数字化素养培训，

GAI 的杠杆效应在于：当整个团队、整个部门、整个组织都掌握了最低限度的 GAI 素养时，协作方式会发生质变。一个人用 GAI 快一倍，十个人用 GAI，整个流程可能快五倍，因为接口变了、交接方式变了、沟通成本变了。而推动这种系统变革的人就是教练。他们的价值不在于自己的 GAI 技能有多高超（虽然通常也不低），而在于他们把 GAI 推广成了企业的能力。他们搭建的不只是模板库，一种把使用 GAI 当作日常的文化，到达第八阶段的人，往往有一种共同的体验：他们不再觉得自己是在使用一项新技术，而是在参与一场社会变革。他们清楚，技术本身是中性的关键在于人们如何加以利用。

扪心自问： “你是否正在主动传播、培训或帮助他人更快地采纳 GAI？”

如果是，那么你已经完成了八阶段的旅程，也有了一个新的起点，当你指导别人的时候，你会从更广的视角重新看到 GAI。

中国 GAI 教练们正在以极具本土特色的方式塑造着 GAI 的普及路径。中国传统职场文化中的师徒制在 GAI 时代找到了新的形式，不是单向的“传帮带”，而是互带：徒弟教师傅怎么用提示词，师傅教徒弟怎么判断什么质量的输出才够好，这种跨代际的知识交换在中国团队中形成了一种独特的活力。

有的企业甚至设立了数字素养大使这类正式岗位。这些教练们自发形成了跨公司的社区，在知识星球、飞书群、微信群里分享经验、交换模板、讨论伦理边界，他们是 GAI 在中国职场落地的真正毛细血管。在中国语境中，教练阶段的中国实践者更愿意称自己为“先行者”或“搭桥人”。他们的目标不是说服

所有人立刻变成 GAI 专家，而是搭一座桥，让那些还在第三阶段徘徊的人能够以一个不那么痛苦的方式穿过“那堵墙”。

作者的自白：一个跨越半个世纪的理想

写下这八个阶段，于我而言，不只是整理一套观察框架。它是一段跨越了近六十年的个人旅程。1966 年，我从台湾中兴大学应数本科毕业，67 年服完兵役，幸运地获得了全额奖学金，飞越太平洋，进入麻省理工学院（MIT）攻读博士学位。那时我选择了计算机科学，一个在当时许多人眼中尚属“新兴的古怪领域”的学科。



一个 1967 年的概念

我加入了 MIT 的 Project MAC（计算机分时系统项目），那是当时全球最前沿的计算研究社群之一，一群疯子在那里畅想“人人皆可计算”的未来。我写的第一篇学期论文，探讨的是一个如今听起来理所当然、但在当时近乎科幻的概念：用自然语言与计算机交流。

是的，你没看错，1967 年，我在试图论证“人应该能用自己的母语跟机器说话”。教授给我的文章打了个 B，说它很有趣。

那是一个连个人电脑都尚未诞生的年代。计算机是庞然大物，由打孔卡和命令行驱动。“自然语言界面”这个概念，在当时就像一个物理学研究生说“我想造一台永动机”一样，被大多数人当作天真甚至可笑的空想。但我放不下它，不是因为我有先见之明，而是因为我隐隐觉得，人机之间隔着的那些“代码”和“命令”，不应该是一道墙，而应该是一座桥。

之后的几十年里，计算技术经历了令人目眩的飞跃，个人电脑、互联网、智能手机、云计算、深度学习……我目睹了一次又一次的“不可能”变成“理所当然”。但在所有这些巨变中，那个 1967 年的念头始终像一颗未发芽的种子，蛰伏在技术史的土壤之下。

直到 2022 年底，ChatGPT 横空出世。我与许多人共同的渴望，终于变成了现实。自然语言交流不再是科幻，而是全球数十亿人每天都在做的事。

但这条弧线也让我看清了一件事：技术的“必然性”只有回头看时才清晰。身处其中时，每一步都充满了不确定性、反对声和失败的尝试。今天我们把 GAI 当作“基础设施”来谈论，但六十年前，连让计算机理解一句简单的话都没门儿。

我想我们何不给 GAI 起个通俗易懂、朗朗上口的中文名称呢？在此，让我提出一个构想，抛砖引玉，敬请赐教：

生成式人工智能，字「爱言」，号「答应宝」
爱言——爱 (AI) 能言善道

答应宝——随时随地答應的宝贝

也许你不太喜欢愛言，那就直接用答应宝好了。

计算思维的视角

其实，作者向来都非常推崇**计算思维**（Computational Thinking, CT），并出版过相关著作。基于计算机科学的理念和实践，计算思维提供了一套强而有力的思维方式，可用于解决问题、设计流程、排查故障以及预防事故等许多方面。计算思维完全可以——也应当——应用于日常事务与决策之中，从而帮助我们提高工作效率并避免不必要的失误。

计算思维主张“快速高效”、“实事求是”、“审慎客观”、“不偏不倚”。在使用 GAI 时，尤其要注意逻辑与推理，步骤和细节，不能盲目相信 GAI 生成的答案或结果。我们应该消化这些信息，确保真正理解，和全面分析，例如“**时事人地物**”各方面。这样我们就能更加仔细地审视得到的材料。凡是要点还要通过多个独立渠道加以核实。要检查 GAI 也要检讨自己。

无论处在哪个 GAI 阶段，我们都应该随时记住“**GAI 可能会出错**”。计算思维者自然会听从这类建议。正因如此，他们尤其能把 GAI 用得更加安全有效。

结语

无论你在上述八个阶段中处于哪个位置，请记住一件事：进程不是一场竞赛，而是一个过程。

目标从来不是“成为 GAI 专家”或“学会写提示词”——那些只是手段。真正的目标是完成一次视角的转换：从把 GAI 看作威胁，转变为把它看作基础设施。知道小心使用，趋利避害。其实，**GAI 也赋予了我们新的自由：一种满足好奇心与求知欲的自由。**

当你完成这个转换之后，你会发现技术本身不再让你焦虑。因为你已经明白了它的边界，也明白了自己的不可替代之处。职业焦虑不再是“它会不会取代我”，而是“我能成为什么样的人”。你不再害怕 GAI 会统治世界或奴役人类，因为你已经亲手使用过它，知道它是好工具，但没有意志。

对大多数人而言，仅仅偶尔接触答应宝并无不妥。随着这场变革的推进，各类主体和服务提供方将会利用它，而大家都会因此享受到更优质且方便的服务及产品。

而有趣的是，恰恰当你不再恐惧它、不再神化它、不再抗拒它的时候，你才能开始真正地使用它去解决那些我们人类自己制造出来的、长期无法解决的难题，无论是气候变化、教育资源不均、医疗资源短缺，还是复杂系统中的决策优化。

答应宝不会自动解决这些问题。它只是一面镜子，映照出人类的意图和选择。但如果你愿意，它可以成为你手中最灵活的那把扳手，去拧紧那些我们一直拧不动的螺丝。这是新的开始。答应宝加油，中国人加油，全世界加油！