

第二章 一次成功由谁承载

《越顺利，越没有你的位置》· 王德生 + Claude 编 · 本章在全书中为第 49–82 页 · ISBN 979-8-90690-018-0

本章原题《一次成功由谁承载？》，作者 阳涌，原文见 <https://sdeuniverses.com/students/yang-yong/bearerless-success/>

摘要

生成式人工智能正在制造一个经验悖论：它一方面能够迅速压缩不同劳动者、学生和专业人员之间的即时绩效差距，另一方面又可能扩大他们在脱离当前模型、提示链、操作者与制度环境之后完成任务的能力差距。现有研究分别把这一现象理解为收入与生产率分化、人力资本与学习分化，或人工智能接入和人机协作条件分化，却共同默认：只要一次任务被成功完成，这项成功所体现的能力必然已经归属于某个可识别的主体一个人、组织或人机团队。本章推翻这一前提，提出“无着成功”与“能力承载极化”两个概念。无着成功是指：一项人工智能中介任务虽然达到了质量标准，但在模型撤除、版本更换、操作者替换或情境迁移后，不存在任何能够稳定复现、解释、修复并对其负责的承载者。能力承载极化则指，不同个人与组织把人工智能辅助成功转化为稳定能力承载者的概率持续分化。由此，AI 时代的两极分化不是现期结果差距，不是个人技能存量差距，也不是工具接入差距，而是成功能否获得稳定归宿的分化。本章提出一张“现期绩效 × 承载稳定性”二维辨别格，构造能力承载率、无着率和承载迁移矩阵，给出八项可证伪命题、五类研究设计以及制度伦理约束。核心判据是：撤去当前模型版本、提示链与原操作者之后，同类任务由哪个登记主体重做，质量下降多少？

关键词：人工智能；两极分化；无着成功；能力承载极化；人力资本；人机协作

Who Bears a Success? Capability-Bearing Polarization in the Age of AI A Theory of Bearerless Success across Economics, Psychology, and Computer Science

Abstract

Generative artificial intelligence creates a paradox that prevailing accounts of inequality cannot fully explain. AI systems may compress immediate performance gaps between workers, students, and professionals while simultaneously widening differences in their ability to reproduce, transfer, repair, and take responsibility for the same work after the original model, prompt chain, operator, or institutional setting changes. Economics commonly observes productivity and wage outcomes; psychology examines the developmental path through which knowledge and judgment are acquired; computer science and human-AI interaction research evaluate the quality of the surrounding collaborative system. Despite their differences, these traditions generally presume that a

successful task performance must already belong to an identifiable capability bearer: a person, an organization, a technical system, or a human–AI team.

This article rejects that presumption and introduces two constructs. Bearerless success denotes an AI- mediated task event that meets an established quality threshold but cannot be reliably reproduced, explained, repaired, and owned by any recognized bearer under predefined perturbations. Capability- bearing polarization denotes a widening difference among individuals and organizations in their probability of converting AI-assisted successes into durable human, organizational, or technical capabilities. The central claim is therefore triply negative: polarization in the AI era is not fundamentally an inequality of current outputs, individual skill stocks, or model access; it is an inequality in whether successful performances acquire a stable bearer.

The article develops a two-dimensional framework combining current performance with bearing stability, formalizes a Capability-Bearing Rate, distinguishes bearerless success from deskilling, cognitive offloading, distributed cognition, automation bias, synthetic competence, and digital inequality, and derives eight falsifiable propositions. It concludes with empirical designs, ethical constraints, and a dated theoretical wager. The zero-modal diagnostic question is: After the current model version, prompt chain, and original operator are removed, which registered actor reproduces the task, and by how much does quality change?

Keywords: artificial intelligence; polarization; bearerless success; capability-bearing polarization; human capital; human–AI interaction

一、一个越来越反常的事实：差距正在缩小，也正在扩大

人工智能进入劳动与学习现场以后，最容易测量的变化通常是“做得更快了”“答得更好了”和“低水平者追上来了”。

在一项涵盖 5,172 名客户服务人员的实地研究中，生成式人工智能辅助使平均生产率提高约 15%，经验较少和技能较低的员工获得了更大的提升；研究还观察到人工智能可能帮助部分员工学习高绩效同事的沟通模式。表面上看，这意味着人工智能具有压缩生产率分布的力量。

在专业写作任务中，Noy 与 Zhang 的实验同样发现，生成式人工智能能够显著减少任务完成时间、提高产出质量，并缩小不同参与者之间的绩效差距。

教育研究则展示了更加尖锐的两面性。Bastani 等人对近千名高中生进行的实验发现，在练习阶段，直接使用通用 GPT 界面的学生表现大幅提高；然而，当人工智能被撤除后，直接获得答案式帮助的学生在独立考试中的成绩反而低于未使用人工智能的对照组。加入教学性护栏、限制直接给出答案后，负面效应基本消失。更值得注意的是，人工智能压缩了练习阶段的成绩差距，却没有同样压缩撤除人工智能后的能力差距。

于是出现了一个不能用“人工智能究竟扩大还是缩小差距”简单回答的事实：

人工智能可能在同一个时点缩小人与人的产出差距，却在另一个时点扩大他们独立继续行动的差距。

一个原来只能完成 60 分任务的人，在人工智能辅助下完成了 90 分；另一个原来能完成 85 分任务的人，也完成了 90 分。从即时结果看，两人的差距几乎消失了。

但七天以后，模型版本更换，原提示模板失效，任务条件略有改变。前者重新跌回 55 分，后者仍能完成 88 分。

那么，第一次共同得到的 90 分究竟意味着什么？

它是两个人的能力已经趋同，还是他们只在一个短暂配置中产生了相同输出？

如果把前一种情形与后一种情形都叫“生产率提高”，我们就把两种具有相反未来的事件放进了同一个统计栏目。

AI 时代真正危险的差距，可能恰恰隐藏在已经被压平的结果里。

二、关于 AI 两极分化的三种主流解释

现有研究关于人工智能与不平等的讨论，大体可以归入三种解释。它们分别来自经济学、心理学与计算机科学，并分别抓住了一部分真实现象。

（一）经济学：谁获得了更多产出与收入

劳动经济学通常从任务配置、劳动生产率、职业结构、工资与就业份额出发讨论技术冲击。

任务型技术变迁理论指出，技术不是简单地“替代某个职业”，而是重新划分人和机器在一组任务中的比较优势。自动化会替代部分既有劳动任务，新任务的产生则可能重新增加劳动需求。人工智能造成何种分配后果，取决于替代效应、生产率效应与新任务效应之间的关系。

在这一传统中，两极分化通常显露为：

- 高技能与低技能劳动者的收益差异；
- 常规任务与非常规任务的职业份额变化；
- 拥有互补资产者与被替代者之间的收入差距；
- 能有效采用人工智能的企业与采用失败企业之间的生产率分化。

这一视角的优势，是它拒绝把技术影响理解成纯粹的个人心理问题。一个劳动者即使十分努力，也可能因为任务被重新组合而失去议价位置；另一个劳动者即使技能没有明显变化，也可能因为其技能与人工智能互补而提高收入。

但经济学的核心观察单位通常是结果：单位时间产出、完成率、质量、工资、就业、利润与市场份额。

只要产出提高，生产率就提高；只要低绩效者追上高绩效者，绩效差距就缩小。

至于这项提高在任务完成以后究竟沉淀到了哪里，往往不是生产率统计首先回答的问题。

（二）心理学：谁真正发生了学习

心理学与学习科学关注的不是某一次答案是否正确，而是行为者在经验之后发生了何种相对持久的改变。

一个学生在提示下解出题目，并不自动意味着他形成了可以迁移的知识结构。一次辅助行为可能促进编码、提取、解释与反思，也可能绕过这些过程，使正确答案在没有学习发生的情况下出现。

认知卸载研究早已指出，人类会把部分记忆、计算与判断要求转移给外部环境。卸载本身并不必然有害：笔记、图表、计算器和搜索引擎都可以释放有限的认知资源。但外部工具究竟是成为认知系统的稳定组成部分，还是使内部能力无法形成，取决于任务、动机、反馈和后续提取要求。

近期关于生成式人工智能与学习的研究进一步表明，辅助期表现与撤除辅助后的表现必须分开测量。直接答案可能提高练习成绩，却削弱独立迁移；带有步骤限制、提示和解释要求的界面则可能降低这种风险。

在这一传统中，两极分化显露为：

- 能把人工智能输出转化为知识结构的人，与只完成任务的人之间的差距；
- 能进行检索、解释、验证和迁移的人，与依赖提示复现的人之间的差距；
- 形成判断力的人，与只形成答案熟悉感的人之间的差距；
- 能在错误和不确定性中学习的人，与越来越回避认知摩擦的人之间的差距。

这一视角关注的是路径：行为者经历了什么加工，知识如何被组织，反馈如何进入下一次判断。

但心理学容易把能力的合法归宿预先限定为个人大脑。只要一个人脱离工具后不能独立完成任务，就可能被判为没有形成能力。

这会遗漏另一种可能：能力虽然没有完整沉淀到个人，却可能稳定地沉淀在组织流程、共享知识库、版本化提示系统、测试套件和可审计技术管线之中。

人类文明的大量能力，本来就不由任何单个人独立持有。

（三）计算机科学：人机系统是否形成有效协作

在人机交互和人工智能辅助决策研究中，关注点既不是单独的人，也不是单独的模型，而是整个协作条件是否支持恰当判断。

Amershi 等人提出的人机交互设计准则强调，人工智能系统需要在初始交互、日常使用、错误发生和长期变化等不同阶段，帮助用户理解系统能力、修正错误并调整信任。

Buçinca、Malaya 与 Gajos 的研究则发现，要求用户先作独立判断、再查看人工智能建议的“认知强制”机制可以减少过度依赖，但这种设计获得了最差的主观评价，而且收益更多出现在认知需求较高的参与者身上。换言之，能够改善判断质量的界面摩擦，也可能降低采用意愿，并对不同使用者产生不平等效果。

近期对知识工作者的研究也发现，生成式人工智能的使用可能改变批判性思考的分布：使用者把精力从直接生成转向验证、整合和监督，而对工具的信心与投入多少批判性思考之间存在系统关系。

在这一传统中，两极分化显露为：

- 能建立恰当信任的人与过度依赖或过度拒绝的人之间的差距；
- 拥有高质量界面、数据、反馈和权限配置的系统与粗糙系统之间的差距；
- 能追踪来源、识别不确定性、修复错误的人机团队与不可审计团队之间的差距；
- 能把模型嵌入稳定工作流的组织与只购买账号的组织之间的差距。

这一视角关注的是条件系统：模型怎样显露不确定性，界面何时制造摩擦，组织如何分配权限，日志是否保存，错误能否被追溯。

它的局限在于，“团队表现良好”可能成为最终结算标准。

只要人机系统总体上做出了正确判断，研究者就可能把它视为成功的协作。至于原有团队配置被拆散以后，这份成功能否被任何主体接住，则不一定进入主要指标。

三、三个学科在同一个案例上给出三种相反判决

设想一家企业部署生成式人工智能客服系统。

一名入职两周的新员工，在人工智能实时建议下，第一次达到资深员工的服务质量；任务完成时间减少，客户满意度提高，投诉率下降。

经济学会说：低经验员工的生产率提升，技能差距正在压缩。

心理学会追问：人工智能撤除以后，他能否独立识别客户意图、处理异常情况并解释判断？如果不能，当前绩效不等于能力形成。

计算机科学则可能说：要求个人脱离系统独立完成，未必是合理标准。只要人机系统能够稳定地产生高质量结果、正确校准信任并在错误时启动人工接管，协作就可以被视为成功。

三种判断都拥有无法被另外两家直接吞并的证据。

经济学不能把真实发生的生产率提高说成不存在。

心理学不能因为当前答案正确，就承认学习必然发生。

计算机科学也不能接受“只有进入个人大脑的能力才是真能力”，因为复杂航空、医疗、软件和供应链系统的能力本来就分布在工具、程序、角色和记录之间。

这不是三种彼此补充的意见，而是三种相互排斥的判决标准：

一次成功究竟由当前结果、能力形成过程，还是协作系统的有效性来裁定？

如果把三种标准并列起来，结论只能是“既要生产率，又要学习，还要设计好系统”。这类结论当然正确，但没有解释最棘手的案例：

- 产出真实提高了；
- 人没有形成独立能力；
- 当前人机系统也真实有效；
- 但系统中的任何一个关键构件发生变化，任务就无人能够接住。

这项成功不是假的。

它也不一定是个人伪装出来的。

问题在于：成功发生了，却没有获得稳定归宿。

四、三种理论共同站在一块没有被检查过的地基上

经济学、心理学和计算机科学虽然争论能力应该如何测量，却共同默认了一条更深的前提：

每一次成功都必然已经归属于某个能力承载者。

对经济学而言，生产率被归属于劳动者、岗位、企业或资本配置。

对心理学而言，学习被归属于个人内部相对持久的认知改变。

对计算机科学而言，协作绩效被归属于人机团队或部署系统。

三家争论的是谁拥有能力，却共同假定能力一定已经被谁拥有。

这一前提在传统工具环境中很少造成严重问题。

锤子不能独立生成施工方案；计算器不会在更新以后改变答案风格；纸质手册不会根据操作者的语气临时重组操作程序。工具参与任务，但任务能力仍然较容易归属于训练有素的人、明确的组织规程或稳定的技术装置。

生成式人工智能改变了这一点。

一次成功可能依赖于以下成分的临时配合：

- 某个操作者对问题的模糊直觉；
- 某一版本模型的隐含行为；
- 一段没有保存的提示链；
- 会话上下文中的偶然信息；
- 外部检索返回的当时结果；
- 操作者对输出进行的少量无记录修订；
- 组织中某位同事在临界处给出的口头提醒；
- 当前平台尚未改变的默认参数。

这些成分共同使任务成功，却没有任何一方掌握完整生成过程。

操作者无法在模型撤除后重做。

组织无法让另一名员工按记录复现。

技术团队无法解释哪一步决定了最终结果。

模型供应商也不对具体业务结果承担责任。

如果几个月后模型版本升级、原操作者离职或任务情境变化，组织只能重新摸索。

此时发生的不是“能力被机器拿走了”。

因为能力可能从未稳定地属于机器。

发生的也不只是“个人没有学会”。

因为当前系统确实完成了过去无法完成的任务。

真正出现的是一种以往统计账本里没有独立位置的事件：

成功存在，但没有稳定的能力承载者。

五、“无着成功”：一种不属于任何单一主体的成功事件

本章把这种事件称为无着成功。

（一）定义

对于一项人工智能中介任务，若它在时点 t 达到预先规定的质量标准，但在一组预先登记的扰动条件下，不存在任何可识别的人、组织或技术系统能够稳定地复现、解释、修复并承担该任务，则该事件构成一次无着成功。

这里的“无着”不是“无人署名”，也不是“无人获得奖励”。

一篇报告可以署在员工名下，却仍然是无着成功。

一段代码可以完整保存在公司仓库里，却仍然是无着成功。

一个决策可以经过主管签字，却仍然是无着成功。

关键不在名义归属，而在扰动之后谁能接住它。

（二）四项最低承载功能

一项成功是否获得稳定承载，至少要考察四项功能：

- 1\、复现：在允许的时间和资源范围内，能否再次完成结构相同而表面不同的任务；
- 2\、解释：能否说明关键输入、判断节点、证据来源与不确定性；
- 3\、修复：当结果出错、条件改变或模型失效时，能否定位并纠正问题；
- 4\、负责：是否存在拥有介入权限、承担后果并可被追问的登记主体。

这四项功能不要求全部集中在一个人身上。

个人可能负责解释与专业判断，组织知识库负责程序复现，测试系统负责错误检测，部门负责人承担最终责任。

稳定承载可以是分布式的。

但“分布”不等于“无主”。

只要组成部分、接口、权限与替代规则能够被识别，系统就存在承载结构。

无着成功指向的恰恰是另一种状态：任务成功依赖一个临时配置，而配置拆开以后，没有任何被承认的结构能够继续承担它。

（三）三类合法承载者

本章承认三种基本承载渠道。

个人承载。原操作者或经过同等训练者，在模型撤除、任务变式或延迟测试中，仍能完成核心判断。

组织承载。原操作者离开以后，替代人员依据组织保存的材料、规范、案例、日志和评审程序，能够重建任务能力。

技术承载。一个经过版本管理、测试、监控和责任配置的技术工作流，在更换操作者后仍能稳定运行，并能够在异常时被解释和修复。

这一定义拒绝把“脱离 AI 独立完成”设为唯一能力标准。

一个组织完全可以把部分能力外置到可靠系统中。

飞机驾驶能力不只存在于飞行员大脑，也存在于仪表、检查单、训练制度、维修体系与空管协作中。

但是，外置必须产生可持续承载，而不是只产生一次漂亮结果。

（四）无着成功不是低质量成功

无着成功最容易被误解为“AI 胡乱生成而人没有发现”。那只是最简单的失败。

无着成功的典型对象恰恰是高质量结果。

它可以通过全部当前审查。

客户满意，文章流畅，代码运行，分析结论正确，任务按时完成。

正因为结果足够好，系统才没有发出警报。

它的风险不在当前失败，而在当前成功掩盖了未来无人接手。

六、AI 时代两极分化的新定义：能力承载极化

“无着成功”描述单个任务事件，“能力承载极化”描述这些事件长期累积后形成的社会结构。

（一）定义

能力承载极化是指：在人工智能广泛参与任务生产的条件下，不同个人、群体与组织把人工智能辅助成功转化为稳定个人能力、组织能力或技术能力的概率持续分化。

因此，AI 时代的两极分化：

- 不是现期产出高低的分化；
- 不是个人技能存量多少的分化；
- 不是人工智能接入与否的分化；
- 而是成功事件能否被稳定承载、继续积累和跨情境调用的分化。

收入、职业、教育和地区差距仍然存在。

本章不是否认这些差距，而是提出：它们越来越可能成为能力承载结构的下游表现。

（二）两极不是“会用 AI”和“不会用 AI”

传统数字鸿沟通常把人分成拥有接入条件者与没有接入条件者。

人工智能素养框架则可能把人分成会提示、会验证、会协作的人与不会使用的人。

能力承载极化划出的两极不同：

一极是可归着增长。每一次人工智能辅助成功，都有一部分被转化为个人判断、组织记忆、技术程序或责任结构；下一次任务因此站在前一次任务的积累之上。

另一极是无着繁荣。任务产量不断增加，短期指标持续改善，但成功依赖的临时配置没有形成稳定承载；模型、人员或情境发生改变后，系统必须重新开始。

两类组织都可能拥有高级模型。

两类员工都可能熟练使用提示词。

两类学校都可能取得漂亮的 AI 辅助成绩。差别在于：

前一次成功是否成为下一次行动可以继承的条件。

（三）为什么这种分化会自我加速

可归着增长具有累积性。

一项任务被解释、保存、测试和制度化之后，后续成员不必从零开始。组织可以比较版本、分析错误、改进流程，并把高水平判断转化为公共资源。

无着繁荣也具有累积性，但累积的是依赖。

每一次快速成功都会减少整理过程、解释依据和训练替代者的短期收益；流程越顺畅，人们越少理由停下来建立承载结构；任务量随后继续增加，使团队更不愿承受记录和学习成本。

于是，短期表现最好的时期，可能正是长期能力最容易被掏空的时期。

这解释了能力承载极化的反常表象：

越成功的人工智能采用，越可能推迟对能力归宿的检查。

七、一张区分四种未来的二维辨别格

只观察当前绩效，无法区分真正的能力增长与无着繁荣。本章引入第二条轴：稳定承载性。

	稳定承载性高	稳定承载性低
当前绩效高	沉淀型增长：任务表现良好，且个人、组织或技术系统能够在扰动后继续承担。	无着繁荣：当前结果优异，但撤换模型、操作者或情境后无人接住。
当前绩效低	潜伏能力：现有结果不佳，但系统保留可解释、可迁移、可修复的能力基础。	显性剥夺：当前表现差，也没有形成可继承的能力承载结构。

表 1 当前绩效 × 稳定承载性的二维辨别格

（一）沉淀型增长

这是多数人工智能政策默认追求的理想状态。

员工借助模型提高效率，同时理解关键判断；组织保存可复用过程；技术系统可追溯、可测试；模型更新后仍能维持工作。

在这一格中，人工智能既提高当前绩效，也扩大未来行动能力。

（二）无着繁荣

这是本章的核心靶格。它具有三个特征。

第一，短期指标表现为改善。产量上升、错误下降、时间缩短、满意度提高。

第二，失效前不产生明显信号。因为当前结果足够好，传统质量控制不会触发。

第三，它会自我加固。指标越好，组织越倾向于扩大部署、减少冗余训练并压缩解释时间。

无着繁荣不是人工智能采用失败。

它是以成功形式发生的能力断裂。

（三）潜伏能力

这一格提醒我们，当前低绩效不等于没有能力。

新手可能工作较慢，却能解释每一步，发现错误并在反馈后迁移。一个组织的流程可能暂时低效，却保存了完整案例、测试和替代机制。

若只按当期产量考核，这类单位可能被无着繁荣者替代。

但在环境变化、模型失效或复杂异常出现时，潜伏能力可能表现出更强韧性。

（四）显性剥夺

这是传统不平等研究最容易识别的一格：没有高质量工具，没有适当训练，没有稳定组织支持，当前结果与未来能力都较弱。

人工智能政策若只把资源投向这一格，而忽视无着繁荣，就会误以为高绩效单位已经不需要能力建设。

真正隐蔽的两极分化，发生在第一行内部。

八、如何测量：从一次任务到一个组织

（一）任务事件

设任务事件 j 的当前质量为 Q_j ，预先规定的合格阈值为 q^* 。

当 $Q_j \geq q^*$ 时，任务被记为一次当前成功。

对每一次当前成功，分别考察三类承载渠道：

- H_j ：个人承载；
- O_j ：组织承载；
- T_j ：技术承载。

每类承载都通过预先登记的扰动测试，而不是通过主观询问判断。可能的扰动包括：

- 1\、撤除当前模型；
- 2\、更换模型或模型版本；
- 3\、移除原提示链；
- 4\、更换操作者；
- 5\、延迟七天或三十天重做；
- 6\、转移到结构相同但表面不同的任务；
- 7\、引入一个需要修复的错误；
- 8\、要求说明证据来源与责任主体。

如果至少一个承载渠道在规定扰动下达到复现、解释、修复和负责的最低标准，则该任务获得了承载。

若 $Q_j \geq q^*$ ，但 H_j 、 O_j 、 T_j 全部未达到标准，则该任务为无着成功。

（二）能力承载率

对于一个人、团队或组织，在观察期内所有合格的人工智能中介任务中，定义：

能力承载率 = 获得至少一种稳定承载的成功任务数 ÷ 全部人工智能中介成功任务数。相应地：

无着率 = 1 - 能力承载率。这一比率与生产率正交。

一个团队可以同时拥有很高生产率和很低能力承载率。

另一个团队可以产出速度较慢，却拥有较高承载率。

这正是新指标存在的必要性：若能力承载率与生产率始终完全同步，它就只是生产率的昂贵替代指标。

（三）承载深度

二元判断仍然过于粗糙。可以进一步把承载深度分为四级：

- 一级：复现承载—能够重做；
- 二级：解释承载—能够说明关键机制和证据；
- 三级：修复承载—能够在错误或条件改变后调整；
- 四级：责任承载—拥有介入权限并承担制度后果。

一个自动化工作流可能具有较高复现承载，却缺乏解释与责任承载。

一个专家可能具有很强解释和修复能力，却因组织没有授权而无法承担最后责任。

因此，能力承载不是单一分数，而是一个承载向量。

（四）承载迁移矩阵

人工智能采用并不只会增加或减少能力，它还会改变能力所在的位置。

一项任务的主要承载者可能发生以下迁移：

- 从个人迁移到组织；
- 从个人迁移到技术系统；
- 从组织迁移到供应商；
- 从技术系统迁移回专业人员；
- 从可识别主体迁移到临时配置；
- 从临时配置重新稳定到组织流程。

前四种未必是损失。第五种才是无着化。

因此，研究人工智能与能力变化，不能只问“人的能力是否下降”，还要问：

能力迁移到了哪里？新的位置能否被识别、维护、替换和追责？

九、从成功到承载：四条稳定化路径

人工智能辅助结果不会自动成为稳定能力。它需要经过至少四种稳定化过程。

（一）认知稳定化

认知稳定化是指，人工智能输出进入人的知识结构、判断习惯与迁移能力。

它不等于阅读过答案，也不等于能够复述答案。关键机制包括：

- 在获得建议前形成自己的初步判断；
- 比较自身判断与模型输出的分歧；

- 解释为什么接受或拒绝模型建议；
- 在相似但不同的任务中重新生成；
- 延迟提取而不是立即模仿；
- 经历能够被定位和纠正的错误。

带有教学护栏的人工智能系统之所以可能优于直接答案系统，不是因为它提供了更多信息，而是因为它保留了能力形成必须经过的部分路径。

然而，认知稳定化不能被理解为“给所有任务增加摩擦”。

认知强制可能减少过度依赖，却降低用户满意度，而且不同认知动机者受益不同。

因此，稳定化设计需要判断：哪些节点必须由人形成可迁移判断，哪些步骤可以可靠外置。

（二）组织稳定化

组织稳定化是指，一次成功被转化为其他成员可以继承的公共能力。它至少需要：

- 保存任务目标、输入和约束；
- 保存关键提示与人工修改；
- 记录采用或拒绝模型建议的理由；
- 建立失败案例而不只保存成功模板；
- 明确任务所有者与替代者；
- 定期由未参与原任务者复做；
- 将隐性经验转化为案例、测试和评审程序。

组织记录的价值不在于“资料越多越好”。

大量未经整理的会话记录不构成组织承载。

只有当下一位行动者能够据此缩短重建时间、识别关键分歧并承担任务时，记录才成为能力结构。

（三）技术稳定化

技术稳定化是指，把临时会话转化为可版本化、可测试、可观察和可回滚的工作流。它包括：

- 固定模型与参数版本；
- 保存提示模板和外部工具配置；
- 建立输入输出测试集；

- 记录数据与证据来源；
- 监控模型漂移；
- 设计失败时的人工接管；
- 保留回滚和复核机制；
- 在更换操作者后验证工作流是否仍可运行。

技术稳定化并不是把任务全部自动化。

相反，一个无法由人解释、无法安全中止的全自动系统，可能拥有复现能力，却没有完整承载能力。

（四）责任稳定化

责任稳定化解决“出了问题谁能在事前介入，而不只是事后背锅”。

名义上签字的人未必是责任承载者。

真正的责任承载至少需要：

- 能够观察任务过程；
- 有权暂停或改变流程；
- 有资源修复问题；
- 会承担可识别的制度后果；
- 其责任范围与控制能力相匹配。

如果一个员工必须对人工智能决策负责，却无权查看模型来源、改变系统或拒绝使用，他不是承载者，而是后果接收者。

能力承载极化因此也是权力极化。

拥有介入权的组织可以把人工智能成功稳定下来；只有使用义务、没有设计权与退出权的群体，则更容易积累无着成功。

十、两极分化如何在微观成功中发生

能力承载极化不是从宏观收入统计突然出现的。它在每一次人工智能辅助任务中发生。

命题一：即时绩效收敛可以与承载能力分化同时发生

当人工智能主要补偿低经验者的即时任务缺口，却没有同步提高其迁移、修复和解释能力时，当前绩效差距缩小，而扰动后的能力差距扩大。

这不是“先缩小、以后又扩大”的时间顺序而已。

在同一时点，只要换一个测量条件，两种方向就可能同时被观察到。

命题二：高产出环境更容易隐藏无着成功

当组织把采用规模、产量、速度和满意度作为主要成功指标时，人工智能辅助任务越顺利，检查能力承载的激励越弱。

因此，无着率未必在低绩效团队中最高，反而可能在短期采用成绩最亮眼的团队中最高。

命题三：模型升级会成为能力归宿的显影剂

若一项能力已稳定沉淀到个人、组织或技术系统，模型更换会造成调整成本，但不会使任务完全失去承载。

若成功依赖当前模型的偶然行为，版本升级会使质量发生断崖式下降，而且团队无法解释下降来自何处。

因此，模型更新、平台故障与供应商切换不是单纯技术噪声，而是识别无着成功的自然实验。

命题四：组织承载可以替代部分个人承载

个人在撤除人工智能后不能独立完成任务，不足以证明能力已经丧失。

若替代员工能够借助组织记录、测试和流程稳定完成任务，该能力已获得组织承载。

因此，把“无辅助个人表现”设为唯一因变量，会系统性低估有效的人机组织学习。

命题五：平等接入不会自动产生平等承载

两所学校即使使用同一模型、拥有相同账号和设备，也可能形成完全不同的能力承载率。

一所学校把人工智能用于答案生成，另一所学校要求先作判断、比较分歧、延迟迁移并保存错误路径。

两者接入条件相同，短期成绩也可能相近，但成功事件的归宿不同。

因此，接入鸿沟缩小以后，能力承载极化仍可继续扩大。

命题六：人工智能摩擦具有分配效应

要求用户先判断、解释理由或验证来源，可能促进能力稳定化，却会增加时间和认知成本。

高认知需求、高先备知识和高自主性使用者可能从摩擦中获益更多；资源紧张者则可能绕过、拒绝或形式化完成这些步骤。认知强制研究已经显示，降低过度依赖的界面并不一定受到用户欢迎，其收益也可能因个体特征而异。

因此，“增加摩擦”不是普遍处方。缺乏补偿和差异化支持的摩擦设计，可能以培养判断力之名扩大新的差距。

命题七：高技能者也可能产生无着成功

无着成功不是低技能者专属风险。

高技能者可能凭借经验判断迅速修改人工智能输出，却没有记录修改依据；最终结果很好，但组织无法知道专家改了什么、为什么改。

Brynjolfsson 等人的研究还发现，高技能员工在使用人工智能建议时可能提高对建议文本的遵循，而部分边际建议并不改善质量；研究也提出，若高水平员工减少原创方案，未来系统可学习的多样性可能受损。

因此，顶尖个体创造的未记录修订，也可能形成组织层面的无着成功。

命题八：承载极化先于收入极化显露

收入和职位变化往往滞后于任务结构变化。

在人工智能采用初期，不同员工可能仍拥有相近职位和工资，但他们获得的任务机会已经不同：

- 一些人持续处理异常、解释冲突和修复错误；
- 另一些人主要提交提示、接收答案并完成格式加工。

前一组不断积累能够接管系统的能力，后一组积累的是依赖当前系统的产出记录。

当组织重组、模型升级或责任事故发生时，这种隐藏差距才转化为职位与收入差距。

因此，能力承载率可能是收入两极分化的领先指标，而不是其同义指标。

十一、五类可实施的研究设计

（一）企业内部多部门自然实验

在同一家企业选择不少于六个使用同一模型、处理相似任务的部门。收集：

- 人工智能辅助期生产率；
- 模型故障或版本更新前后的质量变化；
- 原操作者离开后的替代成本；
- 工作流记录完整度；
- 任务解释与修复时间；
- 个人、组织和技术承载测试结果。

控制任务复杂度、工作总量、员工经验和管理制度后，检验能力承载率能否预测扰动后的恢复速度。

这种设计优于比较两个国家或两家完全不同的公司，因为后者无法排除文化、制度和行业差异。

（二）教育中的双时点随机实验

随机设置四组：

- 1\、无人工智能组；
- 2\、直接答案组；
- 3\、提示与解释组；
- 4\、先独立判断、再使用人工智能并比较分歧组。

分别测量：

- 辅助期任务表现；
- 七天后的无辅助迁移；
- 三十天后的延迟保持；
- 新工具环境中的重新学习速度；
- 学生对自身能力的校准；
- 小组是否形成可共享的解题记录。

该设计不仅比较“是否学会”，还比较能力最终沉淀到了个人、小组流程还是未被任何结构接住。

（三）专业组织中的人员替换测试

从法律、咨询、医疗行政、软件开发或出版团队中抽取已经完成的人工智能中介任务。

在不使用原操作者私人会话记录的条件下，让另一名合格人员依据组织现有材料复做。测量：

- 重建所需时间；
- 与原结果的质量差异；
- 找到关键证据的比例；
- 无法解释的人工修改数量；
- 需要联系原操作者的次数。

这是一种低成本测试，因为多数机构已经保存任务记录、版本信息、工时和人员变动数据。

（四）模型切换实验

选择同一组织中的多个任务单元，在预先登记的窗口内更换模型供应商或主要模型版本。能力承载理论预测：

若组织拥有高承载率，切换会先产生有限调整成本，随后恢复；若组织积累了大量无着成功，切换将产生不成比例的质量下降、返工和责任不明。

关键不是比较哪个模型更强，而是观察：

模型变化以后，组织中是否存在能解释差异并重建流程的主体。

（五）纵向职业发展研究

追踪同一批初级员工三至五年，记录他们在人工智能环境中获得的任务类型。区分：

- 仅提交结果的任务；
- 必须解释理由的任务；
- 必须处理异常的任务；
- 负责模型评估和流程设计的任务；
- 参与错误复盘的任务；
- 能够接替他人工作的任务。

能力承载理论预测，决定未来专业地位的不只是人工智能使用频率，而是员工是否进入了承载形成节点。

相同的工具使用时长，可能对应完全相反的职业结果。

十二、与九种相邻理论的判决性分界

“无着成功”若只是既有概念的换名，就没有理论价值。下面以同一类案例给出可判决的分离线。

认知卸载、技能退化和分布式认知，是本章最重要的传统近邻。

相邻概念	它已经解释了什么	与本章的分离线	判决性观察
技能偏向 技术变迁 与任务极 化	技术改变不同技能 与任务的相对需 求。	它主要预测岗位、任务和收入重新分 布；本章预测相同现期生产率下的扰 动后承载差异。	若即时绩效相同且任务结构不变，但 模型撤换后只有一部分单位崩溃，则 能力承载理论增加了独立解释。
数字鸿沟		本章允许完全相同的工具接入，却产 生不同承载率。	

	不同群体在设备、网络、接入和使用质量上不平等。		若控制模型、账号、设备和使用时长后，承载率仍显著分化，则接入理论不足。
认知卸载	人把记忆或计算要求转移给外部环境。	卸载可以稳定地成为个人—工具系统的一部分；无着成功要求外置后没有稳定承载者。	若替代操作者可依据外部系统稳定复做，则是有效卸载而非无着成功。
技能退化	原有技能因少用、自动化或训练中断而下降。	无着成功不要求行为者曾经拥有该技能，也不要求个人能力下降。	新手从未独立掌握任务，却借助 AI 完成高质量结果，仍可形成无着成功。
自动化偏误	人不恰当地接受自动化建议，尤其在建议错误时。	无着成功可以建立在完全正确的模型输出之上。	若所有输出均正确，但人员和组织在模型切换后无法重建任务，则自动化偏误不能解释。
分布式认知	认知可以跨越人、工具、符号和组织而分布。	本章不否认分布式能力，而是区分可识别、可维护的分布结构与临时配置。	若更换成员后，凭借稳定接口与记录仍能复现，则是分布式承载，不是无着。
恰当依赖	使用者在应接受时接受、应拒绝时拒绝人工智能建议。	一个团队可以在当前任务上依赖校准正确，却没有扰动后的复现与修复能力。	若接受与拒绝决策均正确，但模型版本变化后无人解释失败，则依赖校准不足。
合成能力／合成胜任	AI 使人表现出超过其真实理解水平的专业外观。	本章不以「外观是否虚假」或「个人是否真正理解」为判据，而以任何合法承载者能否在扰动后接住任务为判据。	若个人不理解，但组织流程和技术系统可稳定复现、修复并负责，本章判为已承载，合成能力论仍可能判为虚假。
AI 劳动孤儿化	AI 完成的工作没有被组织账目正确记录或归因。	账目归因与能力承载是两个维度；完整记录的工作仍可能无人复现，未记录的工作也可能已形成稳定能力。	若任务记录完整却在人员替换后崩溃，则不是单纯的劳动记账问题。

表 2 与九种相邻理论的判决性分界（编辑按原稿内容还原行列对应）

2026 年出现的“合成能力”研究，把人工智能制造的专业外观与缺乏内部理解区分开来；另有治理研究把人工智能辅助吞吐量被误认成组织能力称为“synthetic competence”。这些工作已经非常接近本章的问题意识。

但它们仍然倾向于把真实能力与个人理解，或与组织能否调试、扩展工作流联系起来。

无着成功的最小分析单位不是“看起来有能力的人”，而是一次达到质量标准的任务事件；它不预设能力应该落在个人，也不把“人工智能生成”本身视为不真实。只要个人、组织或技术渠道中的任何一个能够稳定承载，任务就不属于无着成功。

因此，最近的概念邻居是“合成能力”，但两者对同一案例可能作出相反判断：

一名医生不具备独立重建模型计算的能力，但医院保存了版本化流程、输入数据、验证规则、替代人员训练和责任制度。合成能力理论可能强调个人理解不足；本章判定该任务已经获得组织—技术承载。反过来：

一名专家深刻理解专业内容，却依靠未保存的私人提示链和临场修订完成任务，离职后无人能够复做。合成能力理论可能仍承认专家真实胜任；本章则把该组织的这次成功判为无着。

两者并非强调程度不同，而是分类标准不同。

十三、十个领域中的理论兑现

（一）学校教育

两名学生都用人工智能写出高质量文章。

第一名能够解释结构选择，在新题目中重新组织论证；第二名只能再次请求模型生成。

两人的当前作品相同，能力归宿不同。

学校若只检查作品质量，会把无着成功认证为学习成果。

（二）教师工作

教师利用人工智能生成完整教案。

如果另一位教师能够依据课程目标、学生反应记录和修改理由继续迭代，教案形成了组织承载。

如果教案只存在于个人会话中，原教师离开后团队无法解释为何如此设计，它便是无着成功。

（三）软件开发

人工智能生成的代码通过所有现有测试，并成功上线。

若团队无法解释关键依赖、无法在需求改变后修改，也没有保存生成环境和评审理由，代码的正确运行可能属于无着繁荣。

这与代码当前是否存在错误无关。

（四）医疗行政

人工智能帮助完成复杂病例编码或资源配置。

如果当前操作者只会接受建议，而组织没有保存依据、复核节点与异常处理程序，那么一次准确决策仍可能没有承载者。

但在诊疗安全场景中，不能为了测试承载率而突然撤除系统。承载测试应优先使用模拟、历史回放和非关键任务。

（五）法律服务

律师借助人工智能检索并生成论证。

若引用准确、结论成立，但没有人能够说明为什么排除了相反判例，任务仍缺乏解释与责任承载。

法律任务的能力不是文字质量，而是面对反例时能够继续辩护和修正。

（六）企业管理

管理者使用人工智能生成战略方案。

一份措辞完美的战略可能获得董事会认可，却没有任何成员掌握其关键假设和反事实条件。

当市场改变时，团队只能重新询问模型，而不能判断原方案的哪一部分已经失效。这是一种高层无着成功。

（七）公共治理

政府部门利用人工智能处理申请、分流案件或生成政策分析。

只记录最终决定而不保存证据、模型版本、人工修改和介入权限，会使决策结果存在却无可追责的能力承载结构。

公共部门中的责任承载不能被“有人签字”替代。

（八）科学研究

人工智能生成假设、代码或分析路径。

研究者可能验证最终数值正确，却无法重建模型如何选择变量、排除方案和生成解释。

结果可重复不等于推理过程可承载；反过来，研究者能解释机制，也不保证计算 workflow 可复现。

科学能力需要个人、组织和技术三种承载的协同。

（九）文化创作

生成式人工智能可能提升个体作品质量，同时降低群体层面的内容多样性。相关实验研究已经提示，个体创造收益与集体新颖性之间可能存在张力。

在本章框架下，问题进一步变成：创作成功是否沉淀为创作者新的审美选择、团队新的工作方法或可持续的技术语言，还是只存在于一次模型配置中。

（十）家庭与个人生活

个人使用人工智能制定旅行、理财、健康或学习计划。

当前方案可能非常合理，但若使用者无法识别关键约束、调整意外情况或承担决定后果，成功规划没有形成完整承载。

AI 时代的依赖不是“凡事问 AI”这么简单。

真正的分界是：询问之后，下一次行动能力留在了哪里。

十四、三个学科反过来迫使理论让步

一个跨学科理论若只借用三个领域的材料，却不允许它们削弱自己的主张，就只是扩张性的命名。能力承载理论必须接受至少三项实质约束。

（一）经济学的约束：不能把所有依赖都称为损失

若没有经济学的反向约束，本章很容易写下：

人工智能辅助下不能独立完成任务，都表明能力发生了空心化。这句话必须删除。

分工、专业化和资本互补本来就意味着个体不需要独立拥有全部能力。现代组织通过把能力稳定地配置在不同岗位、机器和制度中提高生产率。

只要系统能够可靠复现、修复和负责，能力从个人迁移到组织或技术系统可以是增长，而不是损失。

因此，本章不把个人独立性置于最高价值。

它保护的是可持续承载，而不是工具撤除后的个人自足。

（二）心理学的约束：组织承载不能替代所有个人判断

若没有心理学的约束，本章可能写下：

只要组织流程能够复现任务，个人是否理解并不重要。这句话同样必须删除。

在环境变化、异常案例和价值冲突中，组织流程无法穷尽所有条件。若没有成员形成可迁移的概念结构与判断力，流程只能复制既有情境。

因此，在高不确定、高责任和高新颖性任务中，至少一部分关键能力必须形成个人承载。

组织记录可以支持判断，却不能无限替代判断。

（三）计算机科学的约束：承载测试本身可能破坏有效协作

若没有人机交互研究的约束，本章可能主张：

每一次任务都应要求使用者先独立完成，再使用人工智能。

这会把能力承载理论变成低效率的普遍摩擦制度。

对常规、低风险、可逆且拥有成熟测试的任务，强制独立重做可能浪费资源，并降低采用意愿。认知强制机制改善部分判断，却可能带来较差用户体验，并对不同人产生不均等收益。

因此，承载测试应该按任务风险、变化速度和不可逆性分层。

不是每一项任务都要在个人层面证明完整承载。

（四）由三项约束得到的适用边界

能力承载理论最适用于以下任务：

- 条件经常变化；
- 错误后果较高；
- 需要解释和修复；
- 人员与模型可能更替；
- 决策具有不可逆后果；
- 组织声称自己已经形成长期能力。

它不应被机械应用于：

- 一次性娱乐生成；
- 低风险、易验证的日常任务；
- 不需要长期继承的临时表达；
- 已有成熟自动化测试和责任体系的稳定流程；
- 撤除系统本身会制造不必要危险的安全关键场景。

十五、八条证伪条件

本章的核心主张不能靠“未来可能存在风险”维持。以下观察将迫使理论被修改或放弃。

证伪一：承载率没有独立预测力

控制工作量、任务复杂度、员工经验、模型版本与资源投入后，若能力承载率不能预测模型故障、人员替换或情境迁移后的质量变化，则能力承载率只是既有能力指标的冗余代理。

证伪二：即时绩效收敛必然伴随承载收敛

若在多个同质样本中，人工智能每次压缩当前绩效差距时，也同比例压缩扰动后的复现、解释和修复差距，就不存在本章所说的隐藏极化。

证伪三：无着成功无法与合成能力区分

若所有被判为无着成功的任务，都可以仅通过测量个人真实理解与表面表现的差距完整识别，而且组织承载和技术承载不提供任何额外分类，则“无着成功”应被“合成能力”吸收。

证伪四：合法外置与无着化无法实证分离

若替换操作者、模型和情境以后，拥有稳定记录、测试和责任结构的外置系统，与临时会话系统表现没有系统差异，则本章对“分布式承载”和“无着成功”的区分不成立。

证伪五：承载稳定化没有顺序效应

本章预测，组织记录、人员训练、测试体系与责任配置的建立，应当早于扰动后恢复能力的改善。

若恢复能力总是先改善，承载结构才随后补写，或者两者没有稳定先后关系，则本章关于承载形成机制的解释受到反驳。

证伪六：模型切换不显露无着率

若高无着率与低无着率单位在模型更换、供应商迁移和提示失效后的质量下降完全相同，则模型扰动不是有效显影条件，理论必须寻找其他识别机制。

证伪七：个人、组织和技术承载不能互相替代

本章主张三种承载渠道可以在一定范围内替代。如果实证表明，只有个人无辅助能力能够预测长期结果，组织和技术承载始终没有独立价值，则本章必须退回个人能力理论。

反过来，若个人判断在所有任务中都没有独立价值，理论也必须删除对认知稳定化的强调。

证伪八：低成本日志测试不能识别任何差异

许多机构已经拥有系统故障日志、版本记录、人员更替记录、返工时间和质量数据。

若利用这些现成记录，无法发现承载率与恢复时间之间的任何关系，且需要完全依靠主观访谈才能维持理论，本章的可操作性就不足。

最强竞争解释

最强竞争解释是：

本章描述的不就是先备知识、组织成熟度和管理质量的综合结果吗？

为排除这一解释，研究必须在相似任务、相同模型、相近人员经验和相同组织制度内，比较不同工作流的承载形成过程。

只有当承载变量在控制一般能力、组织质量和资源投入后仍预测扰动结果，本章机制才获得支持。

十六、伦理边界：能力承载率不能变成人的评级工具

任何可量化的新指标，都会面临被组织用于考核人的诱惑。

能力承载率尤其危险，因为一个人可能在缺乏权限、训练、记录工具和时间资源的条件下，被迫产生大量无着成功。

把这种结构结果归咎于个人，会再次掩盖真正的承载问题。

因此，应用能力承载率必须遵守三条原则。

第一，指标指向任务位置和工作流，不指向人的本质

低承载率首先意味着某类任务没有形成稳定归宿，不意味着执行者缺乏能力、责任心或人工智能素养。

个人指标只能在充分控制任务类型、权限、训练与组织支持后使用。

第二，不得为了提高指标而在安全关键系统中制造撤除和轮换

医疗、交通、金融结算与公共安全系统不能为了测试“谁能独立完成”而突然撤除关键工具。

应使用模拟环境、历史案例、影子运行、受控切换和非关键任务评估。

第三，任何不利安排必须公开编码规则并提供复核通道

若机构依据承载指标调整岗位、培训、权限或晋升，必须公开：

- 哪些任务被计入；
- 何种扰动被使用；
- 哪些承载渠道被承认；
- 质量阈值如何设定；
- 员工拥有哪些补充证据和申诉权。

更根本的问题是，指标一旦进入考核，组织最便宜的达标方式可能不是建立真实承载，而是补写文件、形式化指定责任人，或者安排容易通过的复现测试。这会制造“承载表演”。

目前看不到一种能够彻底消除这一反噬的制度方案。

因此，能力承载率更适合用于诊断工作流和配置资源，而不适合直接排名个人。

十七、本章自身是否也是一次无着成功

按照本章的判据，这篇文章不能因为结构完整、概念清晰和文献充分，就自动被视为已经形成理论能力。

它同样可能处于“当前绩效高、稳定承载低”的格子。

本章使用人工智能参与资料检索、结构组织、近邻比较和语言生成。若作者不能：

- 解释“无着成功”与相邻概念的核心分离线；
- 在新的行业案例中重新推导二维框架；
- 回应反例并修改适用边界；
- 独立设计承载率的实证测量；
- 说明哪些主张来自证据、哪些属于理论推演；

那么本章本身就是它所批判的对象：一篇达到形式标准却没有稳定理论承载者的论文。

即使作者能够解释，文章也尚未获得组织与学术共同体层面的承载。

它需要经过同行批评、概念替换测试、实证操作化与重复研究。只有当其他研究者能够依据公开定义区分案例、产生反例并修改理论，概念才开始获得公共承载。

因此，发表不是本章能力归宿的终点。

发表只意味着一次理论任务取得了当前成功。

十八、证据分层与本章解释不了的洞

本章的主张可以分为三层。

第一层：无着成功是一类独立存在的任务事件

这是最强、也最容易被推翻的主张。

它要求证明：存在一些成功事件，不能被个人技能、合成能力、认知卸载、分布式认知或组织成熟度完整吸收。

若这一层失败，“无着成功”应被删除。

第二层：二维辨别格具有分析价值

即使“无着成功”不是完全新的存在类别，把当前绩效与稳定承载分开测量，仍可能帮助区分不同人工智能采用结果。

这一层不要求新概念绝对不可替代，只要求二维框架改善预测。

第三层：扰动后的能力不同于辅助期表现

这是最稳的经验主张。

现有劳动和教育研究已经显示，人工智能辅助下的即时表现与撤除或变化条件下的表现不能被自动视为同一指标。不同研究场景的结果并不完全一致：有的工作场景观察到较持久的学习，有的教育场景则发现无护栏使用会损害后续独立表现。

即使本章的新概念最终失败，研究者仍应把辅助期结果与扰动后结果分别报告。

本章尚不能解释至少三个洞。

第一，稳定承载需要持续多久才算“稳定”。七天、半年与五年的结论可能不同。

第二，在高度复杂的系统中，是否可能存在没有任何主体能完整解释，却仍然具有可靠分布式承载的能力。

第三，承载能力本身是否会因过度制度化而失去创造性。一个可复现、可审计的流程可能压缩偏离既有规范的探索。

这些问题不能通过增加定义解决，需要长期经验研究。

十九、结论：AI 时代最深的差距，藏在同样正确的答案之间

关于人工智能与两极分化的争论，通常围绕一个方向性问题展开：

人工智能究竟扩大差距，还是缩小差距？

这个问题把差距理解成一根轴，因而不断得到相互冲突的答案。

人工智能可以缩小即时生产率差距，扩大独立迁移差距；可以降低低技能者进入任务的门槛，又提高未来负责异常情况的能力门槛；可以让更多人获得高质量结果，也可以让更少的人知道结果如何产生、如何修复和由谁承担。

本章提出，必须把问题改写为：

一次人工智能辅助成功，在任务结束以后去了哪里？

它可能进入个人的判断结构。

可能进入组织的公共记忆。

可能进入可维护的技术工作流。

也可能只存在于一个已经消失的人—模型—提示—情境配置中。最后一种不是失败。它曾经真实地成功。

但它没有成为任何主体下一次行动的能力。

由此，AI 时代的两极分化不是收入结果、个人技能或工具接入本身，而是成功能否获得稳定承载的分化。

一端不断把成功变成新的起点。

另一端不断制造没有下一轮继承者的终点。

它们在仪表盘上可能拥有同样的产出、质量和满意度，却生活在完全不同的未来里。

本章最终留下的判据不需要评价“真正的能力”“充分的理解”或“合理的依赖”。

它只问一句可以进入日志、审核与实验的问题：

撤去当前模型版本、提示链与原操作者之后，同类任务由哪个登记主体重做，质量下降多少？

回答不出这个问题时，成功没有消失。消失的是它的归宿。

二十、写死日期的理论赌注

本章作出如下可被公开核验的赌注：

截至 2029 年 12 月 31 日，至少一项覆盖多个同质组织单元的同行评议实证研究，或一套进入机构正式审计流程的人工智能能力标准，将把“模型撤换、操作者替换或情境迁移后的任务复现与修复”列为独立指标，而不再只报告人工智能接入率、使用频率、辅助期生产率或自陈人工智能素养。以下情形不算命中：

- 只在政策文件中笼统提到“保持人的能力”；
- 只要求机构声明保留人工监督；
- 只测量使用者是否信任人工智能；
- 只进行一次无对照的问卷调查；
- 只测试模型准确率而不测试任务承载者；
- 只要求保存提示词，却不进行人员、模型或情境扰动；
- 只发表观点文章而没有可执行测量规则。

若到该日期仍没有出现上述研究或标准，至少说明两种可能：

第一，能力承载不是人工智能治理中的独立问题，现有生产率、学习与可靠性指标已经足够。

第二，本章虽指出一个真实问题，却没有提供足够便宜、可靠和可接受的测量方式。

无论哪一种，本章的核心主张都必须被显著降格。

版本与证据声明

版本：2026 年 8 月 6 日，理论建构初稿。

本章属于概念与理论论文，不把既有单项实验结果直接解释为宏观社会已经发生能力承载极化。文中的宏观机制、能力承载率与承载迁移矩阵均属于待检验理论构造。劳动、教育和人机交互研究只作为部分机制提供经验入口，尚不能单独证明全文命题。

人机分工声明

本章的研究问题、理论方向、核心判断及最终署名责任归作者所有。人工智能参与文献检索、理论对照、结构组织、反例生成与文字表达。作者需对所有引用、概念原创性、事实准确性、伦理边界和最终学术主张承担责任。人工智能不构成本章所称的责任承载者。

附论一·最近邻补切（编辑增补）

原稿第十二节处理了九类相邻理论，但下列五家未出场，其中第一家是「成功没有承载者」这一现象在工程实务里的本名。补切不改变作者的核心判断。

一、关键人依赖（软件与运维实务中的 bus factor / truck factor）

已说到哪一步。软件工程实务里有一个粗糙但极其常用的读数：一个项目里，最少需要多少人同时离开，剩下的人就无法继续维护它。围绕它已有成型的诊断实践——按提交历史计算知识集中度、识别只有一人碰过的模块、把「无人可替」的文件列为风险清单，并以结对、轮岗、文档化与代码评审作为处方。这正是本章所说的无着成功在一个具体行业里的既有命名，而原稿零命中。

分离线。关键人依赖的单位是资产（模块、系统、客户关系），问的是「谁还能维护它」；本章的单位是一次达到质量标准的任务事件，问的是「这次成功之后，复现、解释、修复与负责这四项功能落在哪里」。两条差别是实质的：其一，关键人依赖默认能力曾经存在于某人身上、只是过于集中；本章明确允许一种从未被任何人掌握过的成功（新手借 AI 一次做成，之后谁都不能重做），这是 bus factor 的算法算不出来的格。其二，关键人依赖的处方是分散与备份；本章的第五种迁移（从可识别主体迁到临时配置）指出，即使人数充足、文档齐全，只要成功依赖的是一次性的模型版本、会话上下文与未记录的临场修订，分散人手也不解决问题。

判决性对照。取一批任务，按提交历史与人员冗余算出关键人依赖分数；若在该分数匹配之后，模型替换与操作者替换仍产生不同的复做质量，本章获得独立内容；若关键人依赖分数已完整预测复做结果，本章应并入这一实务传统。

这位祖宗反过来削弱本章的一条。bus factor 有一件本章没有的东西：它可以从现成日志里自动算出来，成本几乎为零。本章的能力承载率要求实施预先登记的扰动测试（撤除、更换模型、更换操作者、延迟重做），成本高出一个量级。若两者在预测扰动后质量下降上表现相当，成本更低者胜出——因此本章第十五节的「证伪八」应当把关键人依赖分数列为必须被超越的基线，而不是只与主观访谈作对比。

二、Nelson & Winter 的组织常规

演化经济学把组织能力安置在「常规」上：常规是组织的记忆，能力靠反复执行而非靠文档保存，且大量常规带有默会成分、无法完整写下来。分离线：常规理论解释能力如何在组织中存续，本章问一次成功是否进入了任何存续结构。判决性对照：若一项任务的复做能力可由该组织既有常规的执行频率完全预测，本章的「组织承载」只是常规强度的代名词；若在常规频率匹配后，保存了拒绝理由与失败案例的单位仍恢复得更好，本章的第九节第二条路径获得独立支持。这位祖宗反过来削弱本章的一条：常规理论会说，本章第九节列出的七项组织稳定化动作（保存目标、保存提示、记录采纳理由、建立失败案例……）大多是显性化动作，而组织能力的主要载体恰恰是默会的、无法显性化的部分；把承载等同于可记录、可交接，可能系统性低估那些靠反复共同执行而非靠文档维持的能力。

三、Walsh & Ungson 的组织记忆

组织记忆被拆成若干存留箱：个体、文化、转换（流程）、结构（角色）、生态（物理布置）与外部档案。本章的三条承载渠道（个人／组织／技术）与之部分重叠，但更粗。建议：本章的「组织承载」至少应拆成流程存留与角色存留两项——原稿第九节第二条路径里的「明确任务所有者与替代者」属角色，「定期由未参与原任务者复做」属流程，二者可以一有一无，合并会掩盖最常见的一种失败（流程齐全但无人被指派）。

四、Nonaka 的知识转化（隐性—显性四模式）

SECI 循环描述隐性知识如何经外化、组合与内化在组织中扩散。本章第九节的四条稳定化路径可以整体读作一次 SECI 的重述。分离线：SECI 关心知识的扩散与创造，本章关心一次成功是否获得可被追问的归属与责任——本章的第四项功能「负责」（拥有介入权限、承担后果、可被追问）在 SECI 里没有对应物，这是本章相对该框架最清晰的剩余。

五、Argyris & Schön 的组织学习

该书正面处理「组织如何学习而不只是个人在组织中学习」，并提出组织的实用理论存在于共享图景与组织物件之中。分离线：组织学习理论关心框架的修改，本章关心一次已完成的成功是否被任何结构接住；一个组织可以完全不修改框架，而把每一次成功都稳稳存进流程（本章判已承载，组织学习判为单环）。编辑注：该作者在本站另一批论文中已被指出同一空缺，宜在下一版正文中一次处理。

附论二·站内划界（编辑增补）

与同批《耦合可逆性极化》——两篇同标「What 第一篇」，须由作者定夺

两篇同日投稿，共享同一个反直觉起点（当下产出差距在收敛、脱离原配置后的能力差距在扩大），共享同一批经验入口（Bastani 的高中数学实验、Brynjolfsson 等的客服现场、Noy 与 Zhang 的写作实验、Doshi 与 Hauser 的创造性研究），并各自画了一张二维格，两张格的靶格是同一格：那篇称「托管型绩效」，本章称「无着繁荣」。

可用的分工线。那篇的对象是可恢复性——条件改变后能否重建问题表征、验证链与修复顺序，测量落在切换损失的分布形状上；本章的对象是归宿——一次成功之后，复现、解释、修复、负责这四项功能是否落到某条可识别的承载渠道上，测量落在承载率与承载迁移矩阵上。本章独有而那篇没有的，是「分布不等于无主」这一区分与六种迁移方向中的第五种（从可识别主体迁到临时配置）；那篇独有而本章没有的，是把「极化」写成必须满足分布判据的经验命题，以及一次真跑过的公开数据复算。按这条线，两篇可改为 What 第一篇与第二篇，本章宜排在后。

一条必须由作者处理的口径冲突。那篇明确规定：只有在稳定产出匹配后二成分模型显著优于一成分、成分间距与人口权重达标、并在至少两类扰动上重复，才可称为「极化」；它据此拒绝把自己的数据称为两极化。按同一判据，本章标题与核心构念中的「能力承载极化」目前没有任何分布证据，应读作待检验的命名而非已成立的结论。本站不代作者改动题名，仅在此标明，并建议两篇统一口径。

与本站《冲突与和平》九篇（同一作者）

那一组的共同判断是「最深的一步在短期指标最好看时完成」，本章的「无着繁荣」在形式上同族。分离线：那一组处理的是资格被取消、转移或新授（谁还有资格触发规则层复核）；本章处理的是能力的归属位置（这次成功落到了谁名下）。可分离的观察是：本章的靶格里没有任何资格被取消，权限、账号与角色表都完好，只是没有任何一方能在扰动后接住任务。

附论三·编辑校订说明

本站在不改动作者判断与论证结构的前提下，做了以下技术性处理：

一、近邻分界表的还原。原稿第十二节的分界表在文档导出时发生单元格错行——「技能偏向技术变迁」一行的分离线与判决性观察被截断并串入下一行，「认知卸载」与「自动化偏误」两行各被拆成两至三行，「AI 劳动孤儿化」一行同样被拆开。本站依据上下文把九类相邻概念逐行合回，只恢复行列对应，未增删任何字词。读者如对某一行的归位有疑问，请以作者原始文档为准。

二、题名中的「极化」目前尚无分布证据，说明见附论二。本站未改动题名与构念名。

三、文献。正文列出 17 条参考文献（含链接），未附 DOI，本站未逐条复核；作者自陈本章为「概念与理论论文」，不把既有单项实验结果直接解释为宏观社会已经发生能力承载极化。第十二节提到的 2026 年「合成能力/synthetic competence」研究在文献表中无对应条目，建议作者补出处。

四、排版。原稿因导出产生的中文字间空格已清理；分节标题、编号列表与二维格按原稿层级还原为网页结构。

五、定位提示。本章自述为 2026 年 8 月 6 日理论建构初稿，五类研究设计均尚未执行，写死日期的赌注兑现期为 2029 年 12 月 31 日。请读者按此定位阅读。

本分册取自《越顺利，越没有你的位置》（德麦国际专著第 84 号）。全书 343 页，含前言、导读、导论、四编十一章、枢纽章、合章、结语、参考书目与三则附录；本章的判据与划界须与枢纽章、合章一并读。全书网页版：sdeuniverses.com/books/m/84/text/