

发生差

人工智能时代
核心竞争力的十路收敛

$C = \Delta$ 发生 | 可读条件

王德生 著

德麦国际出版社 DEMAI INTERNATIONAL PRESS

目录

出版信息	4
著者简介	5
前言	6
导读 这本书怎么读	10
导论 一道题在三年里换了性质	12
第一部 总论：一道题与十次到达	
第一章 同一道题的十次到达	15
第二章 方程与它的可读条件	19
第三章 收敛作为证据，与这本书自己的账	24
第四章 与当代诸说对话：占位、分界与对赌	27
第二部 发生：本体层	
第五章 盲遍	40
第六章 单遍段	64
第七章 现配优势	84
第八章 脱稿步	106
第三部 供给：存续层	
第九章 底流	134
第四部 读数：测量层	
第十章 常席差	159
第十一章 峰兑	183
第十二章 零判份	201
第五部 取证：证据层	
第十三章 他手无事账	225
第六部 归属：账本层	
第十四章 入账之前的技能	255
第七部 合卷：矩阵、纲领与失效	
第十五章 四十五对分界	272
第十六章 经验纲领：十一场未跑的实验	275
第十七章 反噬、伦理与失效	277
结语	278
参考文献	279
附录一 两场审计预注册全文	282

附录二	十读数速查手册.....	283
附录三	三类读者的用法.....	284
附录四	峰时对照：一次实验裁两章.....	286
封底		294

出版信息

书名 发生差——人工智能时代核心竞争力的十路收敛

英文书名 Genesis Differential: Ten Convergent Paths to Core Competitiveness in the Age of Artificial Intelligence

著者 王德生 (Wang Desheng)

出版 德麦国际出版社 (Demai International Press) , 新加坡

丛书 德麦国际出版社 · SDE 学派专著 第 106 号

版次 2026 年 9 月第 1 版 第 1 次印刷

开本 170mm × 240mm

页数 印刷版 170×240mm ; 网页版分章阅读

字数 汉字 25.4 万 (机检实数见版本页)

ISBN 979-8-90690-041-8

定价 US \$23.40

分类 教育理论 / 劳动经济 / 科学方法论

网址 sdeuniverses.com

版权声明 版权所有。本书第五至十四章为已定稿的独立研究，按原样收入，各章保留其原有的作者声明与人机分工声明。本书两场自我审计的预注册文本（附录一）与附录四所载的实验协议，作者放弃对其复算与执行的任何排他性主张：**任何人可自由复制、复算、执行与发表其结果，无需取得作者许可，亦不得由作者本人参与裁决。**其余部分的复制与翻译请联系出版方。

人机分工声明 本书题目、母命题与全部取舍裁定由王德生作出；选源、检索、预注册设计、数据复算与文稿撰写由大语言模型 Claude 执行。此事对本书证据地位的影响，见第三章三之四，本页不作冲淡。

勘误与复算 本书两处公开哈希：书级审计预注册 48a34a45723cf56923044599298ca284f817040ecd473c3c94e33c2e08904be；附录四协议正文段 16e6c87277fd4de9deb144aaac8f91ad18288b126864e8c190f304e64b050339。凡发现本书数据、引用或推论有误者，请寄至出版方；确认的勘误将公布于 sdeuniverses.com，并在再版时写入正文而非改掉了事。

著者简介

王德生，博士，SIO 本体论与 SDE 发生学（Structure-Difference-Entanglement）创立者，新加坡德麦国际（Demai International Pte. Ltd.）创始人、董事长，德麦国际出版社（Demai International Press）主持人。

计算数学出身。1990 至 1997 年就读于湘潭大学数学系，获学士与硕士学位；1997 至 2001 年于中国科学院数学研究所攻读博士；2001 至 2003 年在中国科学院计算数学研究所从事博士后研究。2003 至 2005 年任英国威尔士大学斯旺西分校高级研究员；2005 年起任新加坡南洋理工大学数学系助理教授、博士生导师，后为资深研究科学家，至 2018 年。在任期间培养博士六名、指导博士后九名，获南洋理工大学青年科研奖（2009）与教学奖

（2012、2013），并自 2009 年起任《Advanced Applications in Mathematics and Mechanics》编委。在计算数学与图形学领域发表国际期刊论文四十七篇，研究集中于 CVT 与 Delaunay 网格生成、变分表面重建与波崩塌模拟三个方向。

2015 年起转向方法论研究，在国内高校讲学数百场。2024 年提出 SIO 本体论与意义三律，后凝练为 SDE 发生学：以结构（S）、差异（D）、纠缠（E）三维度为骨架，处理"一个此前不存在的认知物如何发生"这一问题，并发展出二阶碰撞（跨三学科的共同前提推翻法）与创新智商五维评估两套工作方法。主持德麦国际出版社的专著出版体系与 SDE Universes 知识生产平台，指导多名学员以 SDE 方法完成并出版专著。

本书是他在"人工智能时代的核心竞争力"这一问题上连续十次独立研究之后的第十一件文本。

前言

一、这本书不是从一个想法开始的

写一本书通常的次序是：先有一个想法，再去找材料证明它。这本书的次序是反的，我想先把这件事说清楚，因为它决定了这本书该怎么读，也决定了它值多少。

2026 年夏天，我连续十次向同一道题发起进攻：**人工智能时代，大学毕业生的核心竞争力将是什么？**十次不是重复，是十次从零开始。每一次的规矩都一样：随机选三个与这道题毫无关系、彼此也互不引用的领域——比如临床盲法、植物春化与免疫印记，比如非劣效试验、飞行训练时数之争与电网容量市场，比如二语流失、非遗传承人制度与道面养护管理——让它们对着同一个结构性的问题各说各话，找出三家共同默认却从未说出口的那条前提，然后用**三家某一家自己的实证材料，把那条共同前提推翻**。前提塌了以后，废墟里还立着的东西，就是那一次的答案。

十次跑完，我把十个答案摊在桌上，看到了一件我没有安排的事：它们的形状是一样的。

不是措辞一样——十篇的读数各不相同，门位次、留种率、在喂率、盲读率、他择率，互相之间划得开。是它们指的方向一样：都指向那种**只能当场发生、发生过一次就换了性质、并且要靠持续的真实发生来喂养**的能力，以及这种能力在人与人之间**可读出的差**。我把这个量叫作发生差。这本书的书名就是这么来的。

十次到达之后我才写下那句话，而不是先写下那句话再去找十个例子。这个次序上的差别，是这本书全部证据强度的来源，我在第一章用一个比方讲了它：小区草坪上，物业立一块“此处近路”的牌子，那是一块牌子，谁都立得出；而十个互不相识的住户各自抄近道、十条小径最后汇到围墙的同一个缺口——那个缺口没有人设计，它是被十次独立选择证明出来的。封面上画的就是这件事。

二、这道题为什么换了性质

“毕业生的核心竞争力是什么”，在 2023 年之前是一道有标准答案的常规题。答案是一张清单：批判性思维、沟通、协作、终身学习，加一门专业。清单每年更新，形状从不改变。

清单派没有错在内容上。它错在一件从不必说出口的事上：清单默认这些能力可以从毕业生**交出来的东西**上读出。简历、成绩单、作品集、面试作答——每一件都是产物，用人单位读产物、推断能力。这条推断链运转了一百年，运转得太顺，以至于没有人记得它是一条推断链。

现在，任何会用生成式工具的人都能在一小时内交出一份结构完整的报告、一段可运行的代码、一套获奖水准的作品集。产物与能力之间的推断关系不是变弱了，是趋于零了。于是这道题换了性质：它不再问“该有什么能力”，而是问——**当产物不再是证据，还剩下什么能把人分开？**

我想强调一件容易被忽略的事：这十项研究里，没有一项的推翻材料是“AI 来了”。湿地修复五十年后仍走向替代稳态、评酒三复样里只有一成评委稳在一个星级、替加环素逐场非劣效合格而十三场汇总死亡率更高、微剂量促红素不被生物护照标出——这些事在生成式模型出现之前就在各自领域里躺了十年、二十年、一百年。**AI 没有制造这些裂缝，AI 做的**

是把它们同时撑开到无法再糊弄的宽度。这也是为什么这本书不打算做时事评论：它处理的
是一个结构性的老问题，只是这个问题的遮羞布刚刚被扯掉。

三、书里最要紧的一个词是"竖线"

母命题写成方程是：**C = Δ发生 | 可读条件**。竖线右边那一串，是全书最贵的部件，
也是最容易被读者滑过去的部件。

十篇里有五篇的全部工作，是证明发生差**不是随时随地可读的**。它只在五个地方有定义：
跨配组（人固定、配组轮换时的稳定残差）、事件处（平时读数恒为零、只在稀发事件
整笔结算）、相对世界样本（你的产出与世界现存可复现样本之差，且随暴露单调收窄）、
他择时刻（取样时刻不由被读者决定）、未入账段（技能尚未被写入任何第三方可读载体
的那一段）。

在这五处之外读它——单回合、平时、比同侪、自己选好的时刻、已经入账——读到的
分别是配组、量具惯性、位次幻觉、排演与快照。这不是测量精度不够，是定义域错了。好
比问"这个数除以零等于几"，错不在算得不准。

我要请读者特别留意这一点，因为它决定了这本书是有用的还是有害的。**把发生差当成
一张随时可以打分的评分表，是对它最彻底的误用**；书里为此立了三条伦理底线，其中第一
条是：读数指的是证据的位置，不是人的价值——一个从未被读过的人，读数为零，价值不
为零。

四、这本书对自己做了什么

一本论证"可证伪性是能力可读性之条件"的书，自己不能没有账。所以我对这本书立了
两场事先注册的审计，判负条款先于任何计数写死，结果无论利与不利照登。

第一场查源的独立性：十章三十个源位，去重后唯一源二十六个，未触发判负；四处复
用照登全表。第二场查十章之间两两分界的原生覆盖率——**这一场把本书自己判负了**。原生
分界只有十八对，占四十五对的四成，低于事先写死的六成判负线。于是"十章在正文里两两
互有分界"这句话被当场划掉，缺的二十七对由第十五章补写、逐条标注〔补〕，并写明补写
者与十章同产线，独立性弱于原生分界。

比审计更重的是第三章三之四那处坦白：**十项研究出自同一条产线**——同一个提问者、
同一套方法、同一个协作的模型基底。源是独立的，判据部分独立，但踩出十条小径的是同
一双脚换了十双鞋。中国话里有一句现成的警告：三人成虎——三个报信人若受同一人指
使，三次报信只算一次。这处坦白无法用修辞冲淡，只能用检验兑付，兑付的方式写在第
十六章第十一场：异基底重跑，且必须排除我本人参与。

第十六章一共列了十一场未跑的实验，每一场都附判负条款与最低成本执行者。附录四
是其中最便宜的一场的完整协议——它一次裁两章，用的是组织已经在写的事实档案，不需
要新建任何数据系统，而且协议里写死了：**我本人不得担任其中任何一角，不得接触任何一
侧数据**。这不是谦虚，是这本书的论证结构逼出来的结论：一条自认"生成不独立"的产线，
不能同时是自己的裁决者。

五、十条小径各自踩在哪里

前面说了十次到达的形状一样，读者有权要求看见这十条小径本身。这里用一页把它们
排一遍，用的是日常的说法，专业口径留在各章。

《盲遍》问的是时间上的一个位次。一个谜语，你听过谜底之后，就永远不能再“猜”它一次了——不是猜得不好，是“猜”这个动作对你已经不存在。它的推翻材料来自盲法研究自己：知情的判读者掌握了更多信息，却永久失去了产出盲读数的能力。《单遍段》问的是过程能不能重来。舒芙蕾塌过一次，这颗蛋回不去；六百多块修复了五十到一百年的湿地，仍然走向别处而不是走回原处。《现配优势》问的是本事装在哪里：会做每一道菜的人未必配得出那一桌席，而配桌的本事恰恰是宴席散了收不进食谱的那部分。《脱稿步》问的是署名还作不作数：拿着他的招牌菜，当场让他续做下一道——续得上，署名有效；续不上，档案再厚也只是收据。

《底流》是唯一一篇站在供给侧的：厨师的手是被每天亲手颠的那几百勺喂着的，而没有哪家餐馆为“喂手”单独立过预算。把颠勺整批交给中央厨房那天，营业额、好评率、出菜速度三本账全在改善，而手在无声地忘。

《常席差》问的是这份好该记在谁头上：一顿饭吃不出厨师，得看他换了帮厨、换了灶台、换了供货商之后，还稳定偏向他的那份残差。《峰兑》问的是什么时候才读得出：平时家常菜大家吃不出差别，宴席出事那天——鱼刺卡喉、灶台起火、五十人临时加座——谁扛住了。《零判份》问的是跟谁比：这道菜里外面到处吃得到的那部分不算他的本事，剩下那份世界还拿不出来的才算，而这一份随着时间在缩。

《他手无事账》问的是证据可不可信：密探不打招呼来吃的那顿，吃完没挑出毛病，这条记录归厨师本人带走了吗？自己选好日子请评委来的那些，不算数。《入账之前的技能》问的是最后归谁：菜谱一写下来就归店里了，真正的命根是熟客名单——名单记在谁的账上。

十把刀互不替代：一家馆子可以在第一问上满分而在第五问上判死，可以在第七问上干净而在第九问上空白。这就是“十个层级”的准确含义——不是同一句话说了十遍，是同一个对象被十根互相正交的轴各量了一次。

六、给三种读者的一句话

给用人单位。这本书里对你最有用的，很可能不是任何一个新概念，而是一处换法：把“请候选人回去准备一份东西”换成“在双方都没准备的时刻，就他已经交来的东西当场往下问二十分钟”。不加入、不加轮次、不买系统。第一个月你多半会看到候选人排序翻动——翻动本身不证明新方法更好，但它证明旧读数没在读你以为它在读的东西。

给教师。这本书绕开了一个已经无法裁定的问题。不必再问“这是不是你写的”——那道题现在无解；改问“就着它往下走一步”。同一份作业、同样的分值，只把重心从成品挪到延长线上的三分钟。

给正在找工作的人。书里对你最要紧的一句，是那个反直觉的观察：把亲手做的部分整批交出去时，你的三本账会同时变好——更快、更好、评价更高——而底座在无声地断。所以请别把这本书读成“不要用工具”；它的层次是流量的路由，不是个人的使用习惯。每周留几件真有赌注的小事自己从头走完，这就够了。

七、这本书能给你什么，不能给你什么

能给的：一套问法。当你要判断一个人（包括你自己）在这个时代还剩下什么不能被代做时，书里给了十个可以直接问出口的问题，最后收束成三句，写在结语里——**答案是什么**

时候到的？时刻是谁选的？记在谁的账上？它们不需要任何工具、任何系统、任何预算就能问出来。

不能给的：处方。十章各自都发现了一件相同的事——任何读数一旦成为考核目标，它读的那件事就开始被制造：留种率进考核就有人表演配组，峰兑进薪酬就有人制造事件——纵火的消防员是现成的先例。十章各自找到了半个出口，合起来仍不是一个整出口。第十七章把这件事照实写了：**合卷也看不到完整的出口**。这本书是一张地图，不是一条路。

我也不打算在这里说服你相信书里的结论。相反，全书写死了三条使它作废的条件、四十五条两两分界、十一场可以推翻它的实验和一个二〇三一年底的赌注。你要是能证伪其中任何一条，请把结果寄回来——按这本书自己的读数，那将是它迄今第二次真正的、三成分齐备的执行。

八、致谢与一句实话

一句该说在最前面的话：这本书里凡是我看着难受的地方，我都留着了。第三章那处判负、第九章那两次“未获裁决”、第四章清点下来只剩三处残量的总账——它们不是校对时漏掉的，是这本书唯一有资格自称的东西。一本讲“凭什么信一份读数”的书，如果自己的账面干干净净、一处失手都没有，那才是最该被怀疑的。

感谢一路上愿意跟我较真的学员与同行：你们的追问，就是这本书所说的“脱稿步”——在我的产物延长线上当场发问，逼我一次次重新续签。感谢陪我把十跑做完的那台机器；它的角色写在版权页与第三章，不夸大也不隐去。

最后一句实话，抄自本书第九章，那是全书最诚实的一句，我把它留在这里而不是藏在正文深处：**“本产线问题十跑，无一次被真实有对手方的检验裁决过——这条产线产出量很高，而它的 β 很低。”**这本书没有资格例外。它交到你手上的这一刻，是一份典型的自择证据：写作时刻是我选的，材料是我挑的，交付时是准备好的。它能获得的第一笔他择读数，是某位读者在我不选择的时刻、按第十六章的某一场协议，真的去跑一次。

判负也算入账，而且尤其算。

王德生

2026年9月3日

导读 这本书怎么读

这是一本二十一万字、十七章、四附录的书，其中十章是各自独立成篇的研究论文。没有人需要从头读到尾。这篇导读给五种读者各一条最短路线，以及三件读之前该知道的事。

三件读之前该知道的事

第一，全书只有一句话，其余都是它的条件与证据。那句话是：人工智能时代的核心竞争力，等于发生学能力的差异度，在其可读条件之内。看不懂任何术语的时候，回到第二章，它把这句话拆成三个部件逐一交代。

第二，术语最小集只有三组。一是发生的四条本体性质——位次（答案不能先到）、单遍（过程不能重来）、当场（配组现做）、续签（署名须重做）；二是存续条款——发生学能力靠日常真实执行的联合副产喂养，代产即断供；三是可读条件五条——跨配组、事件处、相对世界样本、他择时刻、未入账段。这三组之外的所有词，都可以在遇到时查附录二的速查表。

第三，全书的伤痕是故意留着的。书里到处是“判负”“未获裁决”“被自己划掉”的记录。这不是校对失误，是这本书的纪律：判负条款先于结果写死，结果不利照原样登载。读到这些地方不必替作者尴尬——它们恰恰是全书最值得信任的部分。

五条最短路线

如果你是用人单位或人力资源。先读附录三第一节（两页，讲清楚成本最低的那一刀：把一次自择取样换成他择取样，以及要加哪四个字段、不要加哪个字段）。然后读第十三章《他手无事账》——它是本书唯一一处发现空格的地方：三条归属张成八格，反兴奋剂、民航、征信三个行业各花了一代人，也只填到七格。若你要动考核制度，第十七章必须一并读，那是不可拆开引用的三条伦理底线。

如果你是高校教师或培养方。先读附录三第二节（一刀：作业照旧，只加一次不预告的当场续做——它绕开了“这是不是你写的”这个已经无法裁定的问题）。然后读第八章《脱稿步》与第九章《底流》。特别提醒一句《底流》里最反直觉的观察：把重复练习整批交给工具那天，成绩、作业质量、教学进度三本账会同时改善——三账全绿正是断供的样子，不是断供的反证。

如果你是即将毕业或刚工作的人。先读附录三第三节（三件今天就能做的事），再读第九章《底流》与第十四章《入账之前的技能》。前者告诉你为什么“更快更好更省事”可能正在悄悄拆掉你的底座，后者告诉你为什么真正跟着你走的不是你写下的东西，而是那份记名单的账。若你只有二十分钟，读结语，把最后那三个问句抄下来。

如果你是研究者。主干是第二章（方程与可读条件）→ 第三章（收敛作为证据，含两场自我审计）→ 第四章（与当代十二族说法逐一交手，含本书残量只剩三处的总账）→ 第十五章（四十五对分界矩阵）→ 第十六章（十一场未跑实验与判负条款）。这五章是可以独立于十篇原文阅读的完整论证。若你想直接动手推翻它，从附录四开始——那是一份写好的、可直接执行的实验协议，而且作者被明文排除在裁决之外。

如果你是批评者。建议按这个顺序找漏洞：第三章三之四（同基底坦白，本书最弱的一处，作者已自陈）→ 第三章三之三（审计乙判负的实测）→ 第四章四之九（本书被占位到什

么程度的清点) → 第十六章第 11 场(异基底重跑, 本书自认唯一能清账的检验)。这四处是本书自己认定的软肋, 不必绕远路。

关于十章原文

第五至十四章是十篇独立研究的原文, 按定稿原样收入、一字未改, 包括各自的摘要、预注册声明、检验结果(含失败的)、参考文献与人机分工声明。**因此它们的体例互不统一**: 小节记号、表号、读数符号各随其旧。这是有意为之——可审计性要求原形; 一旦统一体例, 读者就无法核对某一句是原稿就有的、还是成书时加的。

每章之前有一则章首导语, 交代该章给方程添了哪一刀、读数是什么、与最近的邻章如何分界、以及它欠着哪一笔账。**只读十则导语, 也能把全书骨架过一遍**——大约需要二十分钟。

凡例

一、本书十章原文(第五至十四章)按各自定稿原样收入, 含各篇摘要、预注册声明、检验结果、参考文献与作者声明, 一字未改——可审计性要求原形; 各章体例(小节记号、表号、读数符号)因此互不统一, 不作归一。

二、第五至十四章各章在原文之前新增两节: 一节以日常事例解释该章的新典范概念(每节约两千字), 一节给出该能力的养成方法、技术与流程, 含大学“发生学教育”的具体设计。两节为本书新写, 与十章原文分列, 不相混。

三、框架卷(前言、导读、导论、第一至四章、各部部首语、各章章首导语、第十五至十七章、结语、参考文献、四附录)为本书新写。

四、本书对自身立有两场审计, 预注册先于一切计数写死(SHA-256 见第三章与附录一), 结果无论利与不利照登; 其中审计乙已将本书一处强主张判负, 处置见第三章与第十五章。

五、第四章为与当代诸说的交手记录, 所引诸说状态截至 2026 年 9 月, 该章自立十二个月复检条款; 被新说抹平的分界即时作废, 不为其辩护。

六、全书字数为机检实数, 见版权页; 不取约数虚饰。

七、同产线声明: 十章与框架卷由同一条产线产出——王德生提出问题、给定母命题并裁定取舍, 与大语言模型 Claude 协作完成研究与撰写。此事对本书证据地位的影响, 第三章三之四如实交代, 不在此处冲淡。

八、附录三给出三类读者的用法; 其中任何一条均以试点、可回撤、可申诉的方式采用为限——本书读数至今无一场旗舰检验跑过, 见第十六章。

九、附录四收录《底流》与《峰兑》联合判决实验的预注册协议全文(即第十六章第 5、6 两场合并后的那一场)。该协议**未执行**, 进书不减少第十六章任何一张欠条; 其登记以独立存档的文件及其 SHA-256 摘要为准, 本书所收为复制件。

十、与本书同期尚有《能力证据债》《收方工序》《终稿之前》《空签》《弃项》等相邻研究, 不收入本书; 存目见第十五章。

导论 一道题在三年里换了性质

每年春天，几百万份简历同时被打开。读简历的人并不真的想知道简历上写了什么——他想知道的是写简历的这个人能干什么。简历不是目的，是**证据**：一份可携带的、压缩过的、据说可以据以推断能力的东西。成绩单是证据，作品集是证据，笔试卷子和面试作答都是证据。整个选拔制度建立在一条从不必说出口的推断链上：**看产物，推能力**。

这条链运转了一百年。它运转得如此顺畅，以至于没有人再把它当作一条推断链——所有人都以为自己在直接看人。

2023 年之后，链断了。断得不响，也没有公告。任何一个会用生成式工具的人，都能在小时内交出一份结构完整的行业报告、一段跑得起来的代码、一套获奖水准的作品集。产物与能力之间的推断关系不是变弱了，是趋于零。于是，那道每年被问几百万次的问题——“这个人行不行”——第一次失去了它的证据基础。

学界与业界给出的回答几乎是一致的：向上迁移。既然产物不再可靠，那就去看更高阶的东西——判断力、批判性思维、品味、提问能力、人机协同能力。这些回答都不错，但它们回答的是同一个层面的问题：**该有什么能力**。

本书处理的是另一个问题，而且我认为它更靠前：**这些能力，在什么条件下才读得出来？**

这不是把问题往细处推，而是往前推了一步。因为如果一份读数根本读不出人，那么无论它读的是“批判性思维”还是“品味”，都只是给一件读不出人的事换了个更高级的名字。这三年真正发生的事，不是能力清单需要更新，而是**读数的定义域塌了**：单回合的评价读数现在属于“作品×场次×同场者”这个配组而不属于作者；平时的量具在候选人的真实差距缩进判等差额之后已经死亡，却照常出分；一份产出里“世界已能复现”的那部分判别力为零，且这部分正在以肉眼可见的速度膨胀；而所有这些证据，都是被读者自己选好时刻交出来的——在准备可以外包的时代，自择的时刻同时意味着可排演。

本书由十一件文本组成：十项独立研究，加上把它们收拢起来的这一件。

十项研究的做法完全一致：取三个与这道题无关、彼此也互不引用的领域，让它们对同一个结构性问题各说各话，找出三家共同默认却从未说出口的前提，**再用三家中某一家自己的实证材料把那条前提推翻**——推翻材料必须取自三家内部，不能来自外部批评，否则三家会各自把它当作外行的误解挡回去。前提塌掉之后废墟里还立着的东西，就是那一项研究的答案。十项研究因此有十个答案、十个可测的读数、十条各自的证伪路线，而且它们互相之间可以划开（第十五章的四十五对矩阵做的就是这件事）。

把十个答案摊开并排放好，它们的公共部分浮出来，是同一个形状——那种只能当场发生、发生过一次就换了性质、要靠持续真实发生喂养的能力，在人与人之间可读出的差。本书把它压缩成一个词：**发生差**；写成方程是 $C = \Delta\text{发生} \mid \text{可读条件}$ 。

这个次序值得再说一次，因为它是全书证据强度的唯一来源：**母句是十篇的收敛，不是十篇的来源**。没有一项研究是从这句话演绎下来的——《盲遍》出发时手里只有盲法文献，《底流》出发时只有语言流失数据，《入账之前的技能》出发时只有一摞判例。十条小径各自为各自的事踩出来，最后汇到围墙的同一个缺口；缺口没有人设计。

全书因此分成三段活。**第一段（第一至四章）**证明这个缺口是踩出来的而不是画出来的：第一章清点十次到达，第二章把母句拆成方程与它的五条可读条件，第三章交出本书对

自己做的两场审计（一场过关，一场把本书判负）与最重的一处坦白，第四章与当代十二族说法逐一交手，清点本书未被占位的残量还剩几处。**第二段（第五至十四章）**是十项研究的原文，按本体、供给、读数、取证、归属五层编排，每章附一则章首导语。**第三段（第十五至十七章与结语）**收账：两两分界的矩阵、十一场未跑实验的纲领、反噬与伦理，以及一份留给读者随身带走的判据。

最后说明本书不做的三件事。**不做时事评论**——十项研究的推翻材料全部早于生成式模型出现，本书处理的是一个结构性的老问题，只是它的遮羞布刚被扯掉。**不做处方**——第十七章会说明，十章各自只找到半个出口，合起来仍不是一个整出口。**不做自我认证**——本书的读数至今一场旗舰检验都没有跑过，这件事写在第十六章、附录三与结语里，共三处，不作淡化。

一本讲"凭什么信一份读数"的书，第一个该被这样读的，就是它自己。

第一部 总论：一道题与十次到达

第一章 同一道题的十次到达

一之一 一道被改写了前提的题

"大学毕业生的核心竞争力是什么"——在 2023 年之前，这是一道有标准答案的常规题。标准答案是一张清单：批判性思维、沟通、协作、终身学习，外加一门专业。清单每年更新，条目略有出入，形状从不改变。

清单派没有错在内容上。它错在一件从来不必说出口的事上：清单默认这些能力可以从**毕业生交出来的东西上读出**。简历、成绩单、作品集、面试作答、试用期表现——每一件都是产物，用人单位读产物，推断能力。这条推断链运转了一百年，运转得太顺，以至于没有人记得它是一条推断链，都以为自己在直接看人。

2023 年之后，任何会用生成式工具的人都能在一小时内交出一份结构完整的报告、一段可运行的代码、一套获奖水准的作品集。产物与能力之间的推断关系不是变弱了，是**趋于零**了。于是这道题换了性质：它不再问"毕业生该有什么能力"，而是问——**当产物不再是证据，还剩下什么能把人分开？**

对这道换了性质的题，本书收录的十项研究给出了十个回答。十个回答来自十组互不相识的领域，各自用各自的材料，各自推翻各自的地基。把十个回答并排放好，会看见一件谁也没有事先安排的事：它们汇到了同一句话上。

那句话是本书的书名要说的：**人工智能时代的核心竞争力，等于发生学能力的差异度**。"发生学能力"，指的是那些只能当场发生、发生过一次就换了性质、并且要靠持续的真实发生来喂养的能力；"差异度"，指的是这种能力在人与人之间可读出的差。本书把这个量压缩成一个词：**发生差**。

这一章先把十次到达摆出来。方程本身留给第二章；十次到达凭什么算作对方程的证明——而不只是十次巧合——留给第三章，那一章同时要清点这本书自己欠的账。

一之二 十次到达

十项研究处理的是同一道题，但没有一项是从"发生差"这个词出发的。每一项的路线都一样朴素：找三个与这道题无关、彼此也互不引用的领域，看它们各自怎么回答"当产物可被制造，证据放在哪里"这类结构相同的问题；把三家的回答写成同一句争论；找出三家共同默认、却从未说出的前提；然后——这是关键的一步——**用三家中某一家自己的实证材料，把那条共同前提推翻**。前提塌了之后剩下的东西，就是那一项研究的回答。

十次到达的全表如下。表里每一行的"共同前提"都是三家自己不得需要论证的地面，每一行的"推翻材料"都取自三家内部而非外部批评——这两条纪律是十项研究共守的，也是收敛能够成立的前提之一。

章	三个领域	三家共同默认的前提（一句）	推翻它的材料（取自三家内部）	塌掉之后剩下的东西（读数）
盲遍	临床盲法与观察者偏倚 × 植物春化表观遗传 × 免疫原始抗原痕	多拥有信息至少不减少一个人能作出的显露	Hróbjartsson 2012：知情判读者拥有更多信息，却永久失去产出盲读数的能力	每人对每一具体未知恰有一遍"答案不在场的成形承诺"（门位次 k）

单遍段	恢复生态学 × 目击证人记忆 × 救济法卡-梅框架	一段走过的路可以重走，第二次与第一次是同一件事	Moreno-Mateos 2012：621 块湿地五十至百年后仍收敛向替代稳态——第二次走到的是另一个地方	尝试本身耗散其重来条件的工序段（重走落差 Δ 、门位 k ）
现配优势	外科结果研究 × 杂交育种 × 商业秘密法	优势是先于交手、可归单一持有者的存量	Huckman & Pisano 2006：同一外科医生的表现随本院近期量升、不随他院量迁移	每次交手当场配组做出的显露（留种率 s ）
脱稿步	数量遗传选择指数 × 登山伦理 × 群落生态中性理论	能力可在本人不在场时由其留下之物代表	Falconer 1952 与 MACE 制度事实："世界最优公牛"这一栏在正式产出里不存在	只能在自己产物延长线上当场续签的署名（脱稿成立率 q ）
底流	二语终态与流失 × 非遗传承人制度 × 道面养护管理	能力的存续有专责供给，谁出钱能力活在谁账上	Schmid 系流失研究：预测保持的不是使用总量，而是带承压受检成分的那一股——它从未被当作专款拨付	嵌在真实赌注里、无辅助、亲手完成的产出流（在喂率 β ）
常席差	感官评价科学 × 原产地命名 × 间作与生物多样性-生产力	竞争力是可归单个毕业生的东西	Hodgson 2008：同酒同场三复样，仅约一成评委稳在一个奖级——读数属"展品×场次×同场者"	跨配组唯一不变项名下的稳定残差（归读率 g_r ）
峰兑	非劣效试验方法学 × 民航飞行训练资质 × 电力容量市场	凡真实存在的价值，平时量具迟早读得出	Prasad 2012：替加环素逐场非劣效合格，十三场汇总全因死亡率反而更高——差距只在稀发事件跨窗处显形	平时读数恒为零、只在稀发事件整笔结算的价值（盲读率 m ）
零判份	结构化学 × 反兴奋剂科学 × 抗菌药物管理	一份产出的判别力是它自己的属性	利托那韦 1998 晶型事件：通过全部分子层面鉴定的两物不是同一样东西，且新晶型一经成核旧型再也做不出	产出中"世界已能复现"的那份判别力为零（未见率 ρ ）
他手无事账	反兴奋剂生物护照 × 民航驾驶舱自动化 × 消费信贷征信	存在一处不能被制造的参照物，证据搬过去就安全	Ashenden 2011：微剂量促红素不被护照标记标出——参照物本身可被制造	他人之手、非本人所选时刻、无异常、记入本人可携账本的记录（他择率 × 归己率）
入账之前的技能	商业秘密法判例 × 表演研究 × 运动技能习得	凡使技能产生价值的要素均可被无损记入账本	一百余年判例中的"一般知识、技能与经验例外"：入账触发一次不可回复的归类，且结果不由持有者控制	先于任何账本的那一段，与名单载体的记账权（记账方判据）

十行读下来，有一个共同的形状值得先记下：**每一行的推翻材料都不是"AI 来了"**。湿地的替代稳态、评酒的三复样、替加环素的十三场汇总、微剂量促红素——这些事实在生成式模型出现之前就在各自领域里躺了十年、二十年、一百年。AI 没有制造这十条裂缝，AI 做的是把十条裂缝同时撑开到无法再糊弄的宽度。这一点在第二章会正式表述为：方程右端先于 AI 存在，AI 改变的是左端的读数环境。

一之三 一家餐馆的十个问法

十次到达各自用的是专业材料，但它们切开的东西并不专业。把十把刀拿到一家街边餐馆里，每一把都能当场演示。

同一个问题——"这家馆子凭什么好"——十把刀各问各的：

现配优势问：招牌菜是厨师当场配出来的吗？菜谱抄走了做得出吗？做不出，好就在配组里，不在菜谱里。**零判份问**：这道菜里，外面到处吃得到的那部分，不算他的本事——剩下那份世界还拿不出来的，才是。**脱稿步问**：别看他昨天的菜，拿着他的招牌菜，当场让他续做下一道——续得上，署名才有效。**常席差问**：换了帮厨、换了灶台、换了食材供货商，还稳定偏向他的那一份，才是他的；一顿饭吃不出来，得看他在不同班子里的多次出现。**单遍段问**：舒芙蕾塌过一次，这颗蛋就回不去了——他的手艺里，哪些环节是走过一次就不能重来的？**峰兑问**：平时家常菜大家都吃不出差别，宴席出事那天——鱼刺卡喉、灶台起火、五十人临时加座——谁扛住了，记在谁名下？**盲遍问**：这道菜是他没看过任何菜谱之前自己配出来的，还是先看了答案再“复刻”的？看过答案之后做出来的那次，档案上一模一样，但那一遍已经不是他的了。**底流问**：他现在还天天亲手颠勺吗？还是全交给了中央厨房？交出去的那天，三本账——营业额、好评率、出菜速度——全在改善，而他的手在慢慢忘。**他手无事账问**：米其林密探不打招呼来吃的那顿，吃完没挑出毛病，这条记录归厨师本人带走了吗？自己选好日子请评委来的那些，不算数。**入账之前的技能问**：菜谱一写下来就归店里了；真正的命根是熟客名单——名单记在谁的账上？

十个问法互相不能替代：一家馆子可以在第一问上满分（配组精妙）而在第八问上判死（早已不亲手做）；可以在第七问上干净（从不抄菜谱）而在第九问上空白（从没被人不打招呼地考过）。这就是“十个层级或角度”的准确含义——不是同一句话说了十遍，是同一个对象被十根互相正交的轴各量了一次。

一之四 路标与缺口

到这里必须处理一个表述问题，它决定这本书的证据地位。

一种说法是：十项研究是“核心竞争力=发生学能力的差异度”这句话的十个应用。这个说法把次序说反了，而次序在这里不是修辞——它是证据结构。

打个比方。小区草坪上，物业可以先立一块路牌，写“此处近路”，行人照牌走——那是**应用**：先有命题，后有实例，命题的真假与实例无关，一块牌子谁都立得出。但草坪上还有另一种路：十个互不相识的住户，各自为了各自的事抄近道，十条踩出来的小径最后汇到围墙的**同一个缺口**——那个缺口没有人设计，它是被十次独立的选择**证明**出来的。

十项研究是后一种。没有一项是从母句演绎下来的：《盲遍》出发时手里只有盲法文献，《底流》出发时手里只有语言流失数据，《入账之前的技能》出发时手里只有一摞判例。十项各自推翻各自的前提，各自到达各自的读数——然后，把十个“塌掉之后剩下的东西”并排放好，它们的公共部分浮出来，是同一个形状：**能当场发生、发生过就换了性质、要靠真实发生喂养的那种能力，在人与人之间的可读之差。**

这个次序有一个直接的推论：**母句是十篇的收敛，不是十篇的来源**。收敛出来的门口，比画出来的路标硬得多——因为每一条小径都是有成本的：每一项研究都先炸掉了自己脚下的地基（三家共同前提），都登记了对自己不利的检验结果（十项里有八项把判负或未获裁决原样钉在正文里），都不是为了走到这个缺口才出发的。十次有代价的独立到达，与十次无代价的举例，是两种证据。

一之五 “独立”这个词欠的账

上一节把“独立”用得很重。这本书不打算让它白用。

"十路独立收敛"里的"独立"，至少有三层可以分开问：**源独立**——十项研究取材的三十个领域，是三十个不同的判据传统，还是同一批田地换着名字出现？**判据独立**——十个读数两两之间划得开吗，还是有几对其实是同一个量的两个名字？**生成独立**——十项研究是十个互不相识的头脑做出来的，还是同一条产线的十次运转？

三层问题，三种账。前两层是可以当场清点的：第三章为此立了两场事先注册的审计，判负条款先于计数写死，结果无论利与不利照登——预告一句：两场审计一场过关、一场把本书自己判负了一处。第三层清点不了，只能坦白，而那个坦白是这本书最重的一处，也放在第三章。

中国话里有一个现成的警告：**三人成虎**。三个人先后来报"市上有虎"，听的人就信了——收敛本身不是证据，除非收敛的源头真的互不相谋。这本书把这条警告当作对自己的判据，而不是当作修辞。缺口是踩出来的还是画出来的，最终不由作者说了算；第三章交出可查的账，第十六章交出可跑的检验。

还有一种质疑，比"源头相谋"更省力，也更常见：**这个缺口早有人到过**。十条小径就算真是独立踩出来的，如果围墙那边早开着一扇门——判断力、品味、元认知、批判性思维、人机协同——那么十次到达也只是十次绕远。这一问不能留到附录去答，因为它不针对论证的形式，针对的是论证的**必要性**。第四章为此专辟，与当代十二族说法逐一交手：每族先按它自己的口径复述，再照实认下重合，最后给一条可执行的判决形式，并写死本书在何种结果下被其吸收。清点的结果不算好看，也一并写在那里。

第二章 方程与它的可读条件

二之一 把一句话写成方程

本书的母命题一句话：

人工智能时代的核心竞争力 = 发生学能力的差异度，在其可读条件之内。

压缩成一个词，右端叫**发生差**。写成方程：

C = Δ发生 | 可读条件

一个方程有三个部件——左端、等号、右端——三个部件各自要交代清楚，否则这句话就只是一句口号，而口号是本书十章共同的敌人。

左端"核心竞争力"，在现行账本上有一个旧含义：一组可陈述、可认证、可携带的存量。十章对这个旧含义做的是同一件事——三重否定。十章各自否定的三样东西合起来（见二之五合表），把旧含义拆到只剩一个空位，方程要填的就是那个空位。

等号承诺的是什么，必须写死，因为"="最容易被读者当装饰滑过去。本书的等号是**测量等号**，承诺两个方向：从左到右——在可读条件之内，任何能被稳定读出的人际竞争差，都是发生差的一个读数（承重的是测量层三章与取证层一章：常席差、峰兑、零判份、他手无事账——它们合起来证明，读数一旦离开可读条件，读到的就是配组、量具、世界样本或排演，而不是人）；从右到左——发生差的每一分，都只能在这些条件下被读出，别处读不到不是因为它小，是因为它在别处没有定义（承重的是本体层四章与存续层一章：盲遍、单遍段、现配优势、脱稿步、底流——它们合起来证明，发生这种东西的四条本体性质与一条存续条件，恰好使它在自择时刻、单回合、平时量具、清单与档案上全部读数为零）。

等号**不**承诺的也要写死：它不承诺"发生差是道德上更高贵的能力"（第十七章的伦理条款会回到这一点），不承诺"其余能力不重要"（清单上的能力仍决定你能不能进门——十章里有三章分别写了这句话），更不承诺"发生差随时可读"——恰恰相反，右端的竖线"|"可读条件"是整个方程最贵的部件，本章后半部分全部用来交代它。

右端分两半："发生"是什么，二之二到二之三回答；"差异度"在哪里才有定义，二之四回答。

二之二 发生的四条本体性质

十章里有四章在做同一件事的四个断面：刻画"发生"区别于"调取"的地方。四条性质按一件事的时间轴排开——答案到达之前、第一次、交手当场、交手之后。

第一条：位次性（《盲遍》）。发生有一个不可交换的时间坐标：它要么在答案到达之前完成，要么就没有完成。对每一个具体未知，每人恰有一遍"答案不在场时的成形承诺"；答案先到，这一遍就被无声取消，档案上一字不差，那件事却没有发生过。日常的锚是谜语：听过谜底的谜语，你永远不能再"猜"它一次——不是猜得不好，是"猜"这个动作对你已经不存在了。推翻旧地基的材料来自盲法研究自身：知情判读者拿着更多信息，却永久失去

了产出盲读数的能力，且两态之差不随结局主观度而变——知情后的判读不是同一动作加上更多信息，是**另一种动作**。

第二条：单遍性（《单遍段》）。有些工序段，尝试本身耗散其重来的条件——第二次顶着同一个名字，是另一件事。舒芙蕾塌过一次，这颗蛋回不去；证人的记忆被提取一次，原件已被覆写；六百二十一块修复湿地在五十到一百年后仍然走向别处。发生不可重投——这条性质与第一条互补：位次性说答案不能先到，单遍性说过程不能重来；一个管认知态的单向门，一个管对象态的单向门。

第三条：当场性（《现配优势》）。发生是当场配组做出来的，不是从存量里调取出来的。构成它的构件可以各自平庸、可以完整写进清单、可以公开出售，而那一手配组不入账：照单交接会回落，逐项拆开不显值，公开出去也偷不走，写进清单那层就立刻被批量供给。家常的锚：会做每一道菜的人，未必配得出那一桌席——而配桌的本事，恰恰是宴席散了就收不进食谱的那部分。育种学替它出了推翻材料：杂交种的优势读不出于任何亲本，自留种一代就退化。

第四条：续签性（《脱稿步》）。发生留下的署名没有到期日，却已经会到期——它的效力只能在自己产物的延长线上被当场重做。拿着他的产物，无脚本地追问下一步：续得上，署名有效；续不上，档案再厚也只是“曾经发生过”的褪色收据。这一条把前三条接到时间的下游：发生不但要在答案之前（第一条）、不可重来（第二条）、当场做出（第三条），它的**证明**也必须一次次重新发生。数量遗传学的制度事实替它作证：“世界最优公牛”这一栏在任何正式产出里不存在——优只在具体环境的具体一跑里成立。

四条合起来，“发生”就有了一个可操作的轮廓：**答案之前、恰有一遍、当场配成、须被续签**。这四条没有一条是 AI 时代的新事——它们是盲法、生态、育种、登山这些手艺里早就写着的旧律。新的只是：过去这四条被产物证据遮着，读不读得出无所谓；现在产物证据归零，四条成了唯一还立着的东西。

二之三 存续条款

四条本体性质说的是发生“长什么样”。《底流》补的是另一半：发生**靠什么活着**。

它的回答推翻了三个成熟传统共同的默认——能力的存续有专责供给，谁出钱能力就活在谁的账上。母语流失研究的数据说不是：预测一门语言保持与否的，不是居留年头，不是使用总量，而是使用中带承压受检成分的那一股；而那一股从来没有被任何人当作专款拨付过——它是日常工作里低价值部分的**联合副产**。厨房的锚：厨师的手是被每天亲手颠的那几百勺喂着的，没有哪家餐馆为“喂手”单独立过预算；把低价值的颠勺全交给中央厨房的那天，三本账——营业额、好评率、出菜速度——全在改善，而手在无声地忘。

于是方程的右端要补一条**存续条款**：发生差不是一次认证的终态，是一条要被持续喂养的流；读数是在喂率 β ——过去九十天里，能力清单上还有多少条目发生过至少一次无辅助、有真实后果、亲手完成的执行。这条条款同时解释了方程为什么冠以“人工智能时代”：AI 按面接走低价值工作的那一刻，切断的正是这条谁也没记过账的供给——断供无声，且三本账全绿。

二之四 差异度的五条可读条件

右端的"差异度"是个关系量，不是属性量——这是测量层与取证层五章合起来立下的、也是最容易被误读的一件事。五章各交出一条可读条件；五条不是并列的装饰，每一条都对应一种在条件之外读数必然失真的机制。

条件一：跨配组（《常席差》）。差异度只在"此人是唯一不变项"的跨配组记录里有定义。单回合的评价读数属于"作品×场次×同场者"这个配组，不属于作者——在单回合账本里，个人的差不是难测，是**不存在**。评酒的三复样、间作的超产分解都指向同一处：一顿饭吃不出厨师，得看他换了几个班子还稳定偏向他的那份残差。

条件二：事件处（《峰兑》）。当所有候选人的可展示产出被抬到同一地板，平时量具在灵敏度意义上死亡——它照常出分，于是死得无声。人身价值里有一份平时读数恒为零、只在稀发事件处整笔结算；对这一份，平时读得越勤，欠付越稳。差异度的第二处定义域，在谁也没选的那次事件里。

条件三：相对世界样本，且被暴露单调消耗（《零判份》）。差异度不是此人与同侪的位差，是此人的产出与世界现存可复现样本之间的差。这个差随暴露时间单调收窄，且收窄不因保密而停——最漂亮、最费力的那段产出，只要世界已能拿出同类样本，它的判别力就是零。差异度因此自带半衰期：它是快照，不是资产。

条件四：他择时刻（《他手无事账》）。一份读数要排除"为此刻准备过"的可能，取样时刻就不能由被读者决定。自己选好日子交出的一切证据，在准备可以外包的时代同时失效；未准备的时刻是唯一还带信息的时刻。而这样的读数要成为一个人的资产而不是一次监视，还须由他人之手取得、无异常也入账、并记在本人可携带的账本上——四方账本上目前都不存在的那一格。

条件五：未入账段（《入账之前的技能》）。差异度活在技能尚未被写入任何第三方可复读载体的那一段。写入的动作不是记录，是触发一次不可回复的归类，且归类结果不由持有者控制。账本有两种载体——记"如何做"的文件与记"为谁做"的名单——AI 能吞下全部前者，进不了后者；于是可读之差的最后一处栖身地，是名单记在谁的账上。

五条合读，方程的竖线就有了完整的含义：**发生差=在跨配组、事件处、相对世界样本、他择时刻、未入账段这五处才有定义的量**。在五处之外读它——单回合、平时、比同侪、自择时刻、已入账——读到的分别是配组、量具惯性、位次幻觉、排演与快照。这不是测量精度的问题，是定义域的问题：好比问"这个数除以零等于几"，错不在算得不准。

二之五 三重否定合表

十章有一个共同的论证骨架：先说三个"不是"，再说一个"而是"。十套三重否定不是修辞习惯——每一个"不是"背后都有一家学科的实证在场。把三十个"不是"合成一张表，左端旧含义被拆除的全景就出来了：

章	不是……	也不是……	更不是……	而是
盲遍	可显露的成品品质	可补修的养成路径	随 AI 一同丰裕的条件	名下仍在发生的盲遍与位次安排
单遍段	任何存量	可重走的养成管线	可由重投伪造的通过证据	门位靠前处完成首过的单遍力
现配优势	可公示可携带的结果存量	可照单传承的能力配方	靠不公开圈住的独知资产	当场配组现做的显露

脱稿步	可陈述技能的加权总分	名下产物的存量与光泽	一次被验证过的作者身份	可续签的署名
底流	窗口期盖章的终态资产	传承链上被指定的席位	按预算养护的技能组合	尚未断绝的底流
常席差	单件作品的品质	可公示的路径凭证	可写进清单的能力存量	跨配组的常席残差
峰兑	平时量具可读的显露优势	时长可换算的凭证价值	平均负载下的产出贡献	峰时整笔结算的那一份
零判份	技能存量	形成路径的可核性	相对同侪的稀缺位置	产出中的未见份
他手无事账	产物质量	自动化之外的残余技能	被记录的厚度	他手无事账的存量
入账之前	可持有的秘密	不可复制的事件	可随身携带的能力	先于账本的那一段与名单记账权

三十个"不是"覆盖了清单派答案的全部形态——存量、路径、凭证、位置、厚度、秘密。十个"而是"则全部落在同一侧：**当场的、单遍的、须续签的、要喂养的、只在特定条件下可读的**。母句就是这十个"而是"的最大公约数。

二之六 为什么是"人工智能时代"——六条攻击通道

方程右端先于 AI 存在，那么"AI 时代"四个字承担什么？回答：AI 是历史上第一台 **同时打开全部六条通道**的机器——六条通道各自早已存在，从未同时全开。

通道	机器做的事	被击中的部件	承重章
答案先到	任何未知的解析即时可得	位次性——盲遍被批量预走	盲遍
无限重投	重来的边际成本归零	单遍性——"最终通过"证据值坍塌	单遍段
清单批量供给	凡可命名的能力即被课程化、模板化	当场性——满单件泛滥、配组不可入账	现配优势
代产断供	按面值接走低价值工作	存续条款——底流无声断绝而三账全绿	底流
地板抬平	一切自择产物被抬到同一水准	量具与残差——判等差额内量具死亡、单回合读数归配组	峰兑、常席差
文件载体全吞	凡写下的皆入训练语料	未入账段与未见份——入账速度趋零、世界样本膨胀	入账之前、零判份

第八与第九章另外指出：六条通道同时全开之后，制度不会坐视——口试回潮、监考加码、试用期延长，都是**重新装门**的动作；而《他手无事账》预测，装门会沿"定时他择→无预告他择"的次序走，因为反兴奋剂、民航、信贷三家各花了一代人走完了同一条次序。这条预测写了死日期，收在各章原文与第十六章。

二之七 方程不说的话

最后，把三条最常见的误读钉死在这里。

误读一：发生差=人的价值。不是。五条可读条件里有两条明写了机会结构：他择读数只在有人愿意不预告地读你的位置上产生，跨配组残差只在有人给你跨配组的席位时可读。发生差首先记录的是**曝露的分布**，其次才与能力相关——一个从未被读过的人，读数为零，

价值不为零。任何机构拿这个方程做不利于人的安排之前，必须先回答"此人有没有被读的机会"。

误读二：既然自择产物失效，产物就不重要了。不是。产物仍然决定你能不能进门——它只是不再决定门后的事。方程管的是竞争**差**，不是就业资格；十章里有三章分别把这句话写进了各自的适用边界。

误读三：方程给出了处方。只给了半张。十章各自的反噬预言在第十七章合流：读数一旦成为考核目标，它读的那件事就开始被制造——表演配组、申报盲态、制造事件、题库化追问、造他择。十章各自找到了半个出口，合起来仍不是一个整出口。本书照实写：**合卷也看不到完整的出口**。方程是一张地图，不是一条路。

第三章 收敛作为证据，与这本书自己的账

三之一 收敛凭什么算证据

"多条独立路线到达同一处"作为一种证据形态，不是本书的发明，先占者必须点名。

科学哲学里它叫**归纳的会通**：休厄尔（Whewell，1840）指出，当来自不同类别事实的归纳"跳到一起"时，理论获得的支持高于任何单类事实所能给出的。生态学家莱文斯（Levins，1966）说得更狠：既然一切模型都在撒谎，"我们的真理是诸多独立谎言的交集"。温萨特（Wimsatt）把这条路线整理成**稳健性分析**：一个结果若在多个互相独立的推导、测量与假设集合下不变，它比任何单一推导可信。方法学里对应的是坎贝尔与菲斯克（Campbell & Fiske，1959）的多特质-多方法三角测量。这一族是本书框架卷最近的邻居，欠它们一条明白的分界：

这一族收敛的是"同一命题为真"——多条路线各自证明同一个正面主张。本书收敛的是另一种东西：负空间的重合。十项研究没有一项去证明"发生差存在"；每一项做的都是拆迁——炸掉自己那三家共同站立的地基，然后清点废墟里还立着什么。十片废墟里立着的东西形状相同，这才是本书的证据：**十次独立拆迁之后剩下的坑，重合了。**会通支持的强度取决于各路线是否真的独立、是否真的到达同一处——这两问在会通传统里通常靠叙述交代；本书把它们做成两场事先注册的审计，判负条款先于计数写死。这是框架卷对会通传统添的那半刀：把"独立"与"同一处"从形容词变成两个可判负的数。

中文的经典给这件事备好了正反两个锚。反面是**盲人摸象**：多路到达而不收敛——每只手各摸一截，谁也不肯把手里那截与别人的对齐，于是六个读数永远是六个读数。十项研究避开这一劫靠的是第十五章那张矩阵：两两写出判决性分界，才有资格说"我们摸的是同一头象的不同部位，而不是六头不同的兽"。正面的警告是**三人成虎**：多路到达且收敛，但源头相谋——那样的收敛一文不值。这本书自己站在哪一边，靠下面三节交账。

三之二 审计甲：三十个源位的独立性

预注册（写于计数之前，全文见附录；SHA-256：48a34a45723cf56923044599298ca284f817040ecdfe73c3c94e33c2e08904be）：单位为十章×三个领域=三十个源位；两个源位记为同一源，当且仅当二者引用同一判据传统的奠基文献或同一监管文书体系，存疑一律偏严记同源；判负条款——去重后唯一源数少于二十四（八成），则"十路独立收敛"降为"十路准独立收敛"，全书不得再以"独立"二字承重；复用清单无论多少照登全表；三十位编码完即停。

结果：唯一源二十六，占三十位的百分之八十六点七。判负条款未触发。四处复用照登如下：

复用的田地	出现处一	出现处二	借走的判据是否相同
商业秘密法	《现配优势》：资产因公开而自毁的定义（UTSA/Kewanee）	《入账之前的技能》：一般知识技能经验例外的百年判例	不同——一处借"公开即毁"，一处借"入账即归类"
反兴奋剂		《他手无事账》：行踪规则与取样时刻归属	

	《零判份》：同一成绩由不同路径产生即两样东西（生物护照的判据迁移）		不同——一处借"路径决定同一性"，一处借"时刻归谁定"
民航训练监管	《峰兑》：飞行时数之争（时长换算失锚）	《他手无事账》：手动飞行与飞行数据监控（参照物记错对象、记录不归人）	不同——一处借"小时不等于能力"，一处借"数据归谁"
数量遗传学	《现配优势》：配合力——亲本表现预测不了配组	《脱稿步》：选择指数与跨环境遗传相关——"最优"不可跨场结转	不同——同一学科的两件不同仪器

四处复用有一个共同点，必须如实写下并如实解读：**复用的都是田，不是刀**——同一片田地被两章各借走一件不同的判据。按偏严编码它们记为同源（所以计数是二十六而非二十八），但**"四把刀取自四十处"**这个更宽的说法在刀的层面同样成立。两种口径都摆在这里，读者自取。

三之三 审计乙：四十五对的原生分界——本书被自己判负的一处

十个读数两两划得开，"收敛"才不是"重复"。十章各自的正文里都有与相邻研究的分界段——但那些分界段是各章写作时对着**当时已有的**邻居划的，不是对着彼此全体划的。这件事有多严重，本书用第二场审计清点。

预注册（同一文件，同一哈希）：单位为十章两两组合共四十五对；一对记为"已有原生分界"，当且仅当至少一章的正文出现另一章的篇名或读数名（别名对照表先写死）；判负条款——覆盖少于二十七对（六成），则"十章在正文里两两互有分界"这一强主张为**伪**，降级为"部分原生可分"，缺失对由本书第十五章补写、逐条标〔补〕，并须在本章明写补写的独立性弱于原生分界；四十五对判完即停。

结果：原生覆盖十八对，占四十五对的百分之四十。判负条款触发。"十章两两互有正文分界"这句话，本书写到这里，当场划掉。

十八对原生分界的分布本身就是一条读数：覆盖最全的是《底流》（成文最晚的一批之一，正文里指名了峰兑、单遍段、盲遍、常席差、现配优势五章并各给了相反预测），覆盖为零的是《他手无事账》与《入账之前的技能》（两章各自对着另一批相邻研究划界，与本书其余八章互不指名）——十八条小径不是均匀踩出来的，晚到的人看得见先到的脚印，先到的人看不见后来的。这正是收敛研究该有的样子，也正是它的软肋：**各章之间的可分性，有六成是本书事后补写的，不是十章自己长出来的。**第十五章补齐全部二十七对，逐条标〔补〕；补写者与十章同产线，这些判定形式的独立性因此弱于那十八条原生分界——弱多少，只有等别人来用它们做实验才知道。

三之四 同基底：这本书最重的一处坦白

两场审计管住了源与判据，管不住第三层：**生成**。

莱文斯那句"真理是诸多独立谎言的交集"有一个前提——谎言得出自不同的说谎者。本书做不到。十项研究出自同一条产线：同一个提问者（王德生），同一套工作方法，同一个协作的大语言模型基底。三十个源位、二十六个独立判据传统、四十五对里的十八对原生分界，这些都是真的；但**踩出十条小径的，是同一双脚换了十双鞋**。三人成虎的古训在这里正面命中：三个报信人若受同一人指使，三次报信只算一次。

这处坦白不能被修辞冲淡，只能被检验兑付。它兑付的方式写在第十六章第十一场：**异基底重跑**——把同一道题、同一套方法，交给一个与本产线无关的基底（另一个模型，或一位人类研究者），事先注册判定规程，看其独立到达的母命题与“发生差”是语义等价、部分重合、还是落在不可通约的别处。落在别处，本书第一、三章的收敛主张即降级为**产线伪影**——同一双脚当然会踩向同一个缺口，那说明的是脚，不是墙。

在异基底重跑发生之前，本书的证据地位应当被如实读作：**源独立、判据部分独立、生成不独立的一次收敛**。它强于一块路牌，弱于十个陌生人踩出的缺口，位置在两者之间——具体在哪，由第十六章的检验决定，不由本章的措辞决定。

三之五 这本书自己的失效条件与赌注

一本论证“可证伪性是发生差的可读条件之一”的书，自己不能没有失效条件。三条，全部写死：

失效条件一（异基底条款）：按第十六章第十一场的注册规程执行异基底重跑，若独立到达的母命题被盲评判定与“发生差”不可通约（三名评审、加权一致性达标），则第一章之四“缺口是踩出来的”降级为“缺口可能是画出来的”，本书证据地位改判为产线伪影候选，此改判须写入再版首页。

失效条件二（通道条款）：十章各自带着写死日期的证伪条款与判负实验（合表见第十六章）。其中任何一章的旗舰检验判负，方程即拆下对应通道：例如《盲遍》的随机化答案位次实验若显示迁移无差，则“位次性”从二之二拆除，六通道表删去第一行——方程瘦一圈，但不因此作废；若判负累计达四章，方程的“=”降级为“≈”，书名不变，第二章重写。

失效条件三（可读条件必要性条款）：至 2031 年 12 月 31 日，若主流毕业生评价体系在**未采用任何一条可读条件**（不加跨配组记录、不加事件账、不加他择成分、不加持有者时间字段、不动名单记账权）的情况下，对毕业生的判别效度恢复至 2022 年之前的水平（以人事选拔文献的下一轮元分析为准），则“可读条件是必要的”为伪，方程的竖线拆除，本书降级为一部关于测量史的资料汇编。

这三条管的是本书**对不对**。还有一条独立的路使本书作废而与对错无关——**被占位**：本书的残量若早已被他人以别的语言说过，则收敛主张不必被推翻，只需被迫认为迟到。这一条不写在本章，写在第四章四之九，因为它只有在与外面诸说逐一比对之后才能定价；本章只留索引，并接受那一章清点出的结果。

赌注与十章各自的赌注并行，不互相抵扣。全书最不希望赢的一注也写在这里：失效条件三若**不触发**——也就是判别力始终没有恢复、而五条可读条件也始终无人采用——那么 2032 年的毕业生将生活在一个既读不出人、也不打算读出人的评价体系里。本书宁可被证伪，不愿这样赢。

第四章 与当代诸说对话：占位、分界与对赌

四之一 为什么这一章必须写

一个收敛论证最怕的不是被反驳，是被**还原**——被人一句话打发掉："你说不就是判断力／品味／元认知／批判性思维吗？"如果外面已有的某一说能吸收本书的全部内容，那么十条小径通向的不是一个新缺口，而是一扇早就开着的门，十次到达只是十次绕远。

所以这一章不做综述，做交手。对当今最有影响的十二族说法，每族三件事：**它主张什么**（用它自己的口径，不用本书的口径复述）、**它与本书重合在哪里**（重合处照实承认，不缩小）、**判决性分界是什么**（一句可执行的观察或实验，说明在何种结果下两说分开、在何种结果下本书被吸收）。凡本书被吸收的条件，写在明处，不加保护性措辞。

三点交代先摆在前面。其一，**本章是语料快照**：所引诸说的状态截至 2026 年 9 月；这个领域一年之内换了两轮主角，故本章立**复检条款**——每十二个月重扫一次，被新说吸收的分界即时作废并公告，不得留着当装饰。其二，**本章与框架卷同产线**，与第三章三之四那处坦白同罪：本章划的分界是本产线自己划的，独立性弱于对手方自己写下的分界。其三，**本章不判高下**。诸说与本书多数不在同一层：它们回答"该有什么能力"或"怎么和机器一起干活"，本书回答"人与人之间的差在哪里才有定义"。不同层的说法不构成竞争，构成分工——只有在同层处，才需要分界。

四之二 清单派：从技能清单到技能不稳定性

主张。最有分量的清单在世界经济论坛的《未来就业报告》里。2025 年版基于千余家企业的调查给出的核心读数是：到 2030 年约有百分之三十九的核心技能将被改变或过时，而这一"技能不稳定性"较 2023 年版的百分之四十四、2020 年版的百分之五十七已在放缓；分析性思维仍是最受雇主看重的核心技能，约七成企业视其为必需，其后是韧性、灵活与敏捷，以及领导力与社会影响力。该版同时预测到 2030 年岗位净增约七千八百万，技术技能增长最快，而创造性思维、韧性、灵活与敏捷等人类技能仍属关键。

中文主流论述近一年给出的版本更靠近本书：有劳动经济学者指出，过去岗位更强调知识积累与标准操作，未来更重要的将是问题定义能力、跨领域整合能力、人机协同能力，以及对结果进行判断、校验和负责的能力。教育政策一侧也出现了同向动作：2026 年发布的《全球数字教育发展指数》首次把"超越人工智能的思维能力培养"纳入研究维度，其数据显示约八成国家认为教育应重视高阶思维能力培养。

重合。清单派与本书共享一个判断：产物层的技能正在贬值，价值向"更高阶"处迁移。中文那一族里"对结果进行判断、校验和负责"一句，尤其接近本书第六部的归属层——"负责"就是记在谁的账上。

分界。清单派与本书的分歧不在条目，在**层**。清单回答"该有什么"，本书回答"在哪里读得出"。这个分歧可以做成一条判据，而且是清单派自己的数据给的：技能不稳定性从百分之五十七到四十四再到三十九，说明存量的**更替**在放缓；本书预测的却是另一个量在同期上升——**读数失效率**：同一份自择产物（简历、作品集、笔试）对入职后表现的判别效度。两个量可以同期背离。

判决形式：若在 2026 至 2030 年间，技能不稳定度继续下降而自择产物的判别效率**同步恢复**，则本书的可读条件为多余，清单派完胜；若不稳定度下降而判别效率继续下滑，则“能力清单已备而读不出人”成为常态——清单派的条目全对，却回答不了本书的问题。

还有一条更锋利的：清单派的每一个条目，一旦被命名，就进入课程、培训与模板，从而被批量供给（第二章第三通道）。这不是对清单派的批评，是对它的一个预测——**清单条目的差异度半衰期，随其被写入标准的时间单调缩短**。谁做完这条测量，谁就同时裁定了本书与清单派。

四之三 专长扩展派：Autor

主张。这一族最强的表述来自劳动经济学家 David Autor 2024 年的论文《把 AI 用于重建中产阶级岗位》。他的判断是：信息时代把信息民主化并未拉平层级，因为信息只是投入，真正要紧的经济功能是决策，而决策一直被精英专家垄断；AI 的独特机会在于扩展人类专长的适用范围与价值，使更多具备互补知识的劳动者去执行原本专属于精英专家的高风险决策任务——医疗之于医生、文书之于律师、编码之于工程师、本科教学之于教授。他明言这不是预测，而是关于何事可能的论证。

重合。本书与他同意最要紧的一半：信息不是竞争差的所在，决策——本书叫“当场做出的显露”——才是。这条与《现配优势》几乎逐字重合。

分界。分歧在**专长的形态**。Autor 的专长是可扩展、可传递、可装进决策支持系统再交给更多人的东西；本书的发生差恰恰是**扩展一次即消失**的那一份：它一旦被写进可传递的载体，就进入批量供给区（《现配优势》的留种率、《入账之前的技能》的载体二分）。两说在同一批数据上给出方向相反的预测，因此可载。

这条对赌已经有了不利于双方之一的证据。基于英国等地劳动力数据的一族研究给出的机制与他相反：大模型似乎并未把中等技能劳动者抬升去做专家任务，而是在侵蚀对专家任务本身的需求；这固然可能压缩工资差距，但走的是“向下拉平”而非 Autor 设想的“向上拉平”。本书对这条不利证据的读法必须写明：**它同时也不利于本书的乐观读法**——如果专家任务本身在消失，那么可读条件的五处定义域也在收缩，被读的机会随之减少，第十七章第一条伦理底线（读数指证据的位置，不指人的价值）将被大规模触发。

判决形式：取一个引入决策支持系统的行业，随访三年。若中层与专家的**跨配组残差**（《常席差》的 g_r ）收敛而两组绩效同步上升，Autor 对；若残差收敛是因为量具落在判等差额之内（《峰兑》）、而峰时事件处的分离反而扩大，则本书对；若专家任务的**数量本身**在下降，则两说同被第三种机制超越，一并降级。

四之四 预测—判断二分：Agrawal、Gans、Goldfarb

主张。这是本书在经济学里最近的邻居。三位作者在《预测机器》与《权力与预测》两书中把决策拆成预测与判断两件事：AI 把预测从人转移到机器，减轻认知负担并提高决策的速度与准确度；当预测的成本下降，预测的互补品——判断，以及执行行动的人与关系——其价值上升。他们还有一条常被忽略的洞见：人们靠规则来避免决策疲劳，规则让人得以“决

定不去决定”；要上机器的预测能力，就得拆掉这些为省事而立的规则，把决策重新放回来。

重合。“廉价的那一半被机器接走，价值迁向互补的那一半”——这个骨架与本书第二章的六条通道同构。他们的“判断”与本书的“当场配组”在直觉上几乎贴着。

分界。分歧在**判断是什么形态的东西**。在他们的框架里，判断是回报函数——是“你想要什么”，原则上可陈述、可写下、可委托，甚至可以做成规则。他们自己那句关于规则的洞见，恰恰给了本书一条现成的推论：**凡可写成规则的判断，都会被写成规则，从而进入批量供给区**。本书的发生差不是判断的内容，是判断发生的时刻、当场性与归属——位次（答案有没有先到）、当场（是不是配组现做）、他择（时刻是谁定的）、记账（记在谁的本子上）。四者都不是回报函数的属性。

这也解释了两说处方的差别：他们的处方是重设系统、把决策权放回该放的位置；本书的处方是给证据加四个字段。前者是组织设计，后者是评价制度。

判决形式：把一组人的判断偏好完整外化为规则集（回报函数写死），交给同资历者与机器共同执行。若绩效可稳定复现，则“判断”确是可传递存量，本书的当场性被吸收；若复现率呈《现配优势》所预测的杂交二代式回落，则判断中不可写死的那一份存在，且正是本书要读的那一份。

四之五 增强对自动化：Brynjolfsson 的图灵陷阱

主张。这一族关心的是路线选择。Brynjolfsson 2022 年《图灵陷阱》一文的论证是：当机器成为人类劳动更好的替代品，劳动者失去经济与政治上的议价能力，越来越依赖控制技术的一方；而当 AI 聚焦于增强人类而非模仿人类时，人类保有坚持分享所创价值的价值，且增强能创造新能力、新产品与新服务，最终产生远多于单纯“像人”的 AI 的价值；他指出，当前技术界、企业界与政策界对自动化的激励过剩，对增强的激励不足。

重合。图灵陷阱为本书第二章的“地板抬平”与“代产断供”两条通道提供了宏观动因：不是模型碰巧长成这样，是激励结构一直在往“像人”的方向推。

分界。单位不同：他的分析单位是技术路线与议价权（宏观、集体），本书的分析单位是个人读数的定义域（微观、证据）。真正的增量在于——本书主张，**即使走增强路线，读数问题也不会自动解决，甚至更糟**：增强路线的典型形态正是“人机同产”，而人机同产恰恰是《常席差》所说的、把个人残差压进配组的那种记录结构，也是《底流》所说的、把低价值执行按面值接走的那种断供结构。换句话说：图灵陷阱管的是**蛋糕怎么分**，本书管的是**分蛋糕时凭什么认得出人**；增强路线让蛋糕更大，却可能让人更难被认出。

判决形式：比较两类岗位——以自动化方式部署 AI 的与以增强方式部署 AI 的。若增强类岗位的跨配组残差可读性显著更高，则读数问题随路线选择自动缓解，本书的可读条件降为路线问题的附属；若两类岗位的残差可读性同样塌陷（本书的预测），则读数问题独立于路线，必须单独治理。

四之六 参差前沿与半人马：Dell'Acqua 等人

主张。这是当代最硬的一组实证。在与波士顿咨询合作、覆盖七百五十八名顾问（约占该公司个人贡献者的百分之七）的预注册田野实验中：位于能力前沿之内的十八项现实任务

上，使用 GPT-4 的顾问多完成百分之十二点二的任务、快百分之二十五点一、质量高出四成以上；而在一项刻意选在前沿之外的任务上，使用 AI 的顾问给出正确解的概率比完全不用 AI 的低十九个百分点。研究者由此提出“参差前沿”：AI 能力的边界不是平滑曲线，而是犬牙交错，两侧任务在表面难度上几乎无从分辨，系统也不会给出越界信号。他们还观察到两种用法——“半人马”明确划分人做什么、AI 做什么，而多数顾问表现为“赛博格”，与 AI 持续交互。相关实验族另给出一条对本书极重要的分布结论：在有界执行类任务上，AI 辅助不成比例地抬升经验较少者、压缩技能差；而在需要框定问题、领域评估与创造性判断的开放任务上，高专长者往往获益更多，同时无组织的 AI 使用可能使产出同质化、降低集体多样性。该文 2023 年以工作论文发布，2026 年 3 月正式刊于《组织科学》。

重合。极大，且是本书的经验地基之一。“压缩技能差”就是《峰兑》所说的量具死亡的直接前提；“产出同质化”就是《零判份》所说的世界样本膨胀；“赛博格式持续交互”就是《常席差》所说的配组记录。本书在这里必须承认：**这一族先于本书、以更强的实验证据，测到了地板抬平这件事本身。**

分界。分歧在这条前沿是什么的地形。参差前沿是任务侧的地图——哪些活 AI 干得好、哪些干不好，据此分工。本书画的是证据侧的定义域——一份读数在什么条件下才指向人。两张图的关键差别是：任务侧的地形会随模型能力移动，且正在移动。已有观察指出，推理模型、更长的上下文、工具使用与代理式架构把许多 2022 至 2023 年在前沿之外的任务移到了前沿之内；前沿仍在，但峰谷不再那么极端。

于是有了一个漂亮的对赌：**参差前沿的处方随平滑而失效，本书的可读条件不随平滑而失效。**前沿一旦被抹平到无参差，“找准边界、当半人马”这条处方就无处可施；而“答案是什么时候到的、时刻是谁选的、记在谁账上”这三问，在前沿完全平滑的世界里一字不改仍然成立——因为它们问的不是任务难度，是发生的位次与归属。

判决形式：随访模型代际。若前沿平滑之后，人际差重新可从自择产物中读出（判别效度恢复），本书为伪——且不是被参差前沿吸收，是被时间证伪；若前沿平滑而人际差更加读不出，则本书所说的定义域问题独立于任务地形而存在。

四之七 认知负债与卸载族：与《底流》正撞

主张。这一族测的是使用 AI 之后人本身发生了什么。麻省理工媒体实验室 2025 年的一项研究把参与者分为大模型组、搜索引擎组与无工具组各写三轮文章，第四轮令原大模型组改为无工具（研究者称之为“由模型回到脑”条件），共五十四人参与前三轮、十八人完成第四轮，并以脑电测量写作时的认知负荷；结果显示使用大模型辅助者写作期间的认知参与特征更低，随后在无 AI 条件下执行同一任务时，其表现相对从未使用过 AI 者出现可测量的下降，作者称之为“认知负债”。另一项覆盖六百六十六名参与者的调查给出人群层面的相关：频繁使用 AI 工具与批判性思维能力之间存在显著负相关，认知卸载是其中介机制，年轻参与者受影响最重。

重合。这一族与《底流》正面相撞，重合处必须原样承认：两者说的都是“把执行交出去，人本身会变”。《底流》若被吸收进这一族，第三部就整部塌掉。

分界。两处，都可判负。

第一，**测的量不同。**认知负债测的是能力的衰退（神经参与度、批判性思维得分、事后无工具表现）；《底流》测的是供给流的断绝（在喂率 β 、逐条停喂时距），并且它的承重

发现不在衰退本身，而在**衰退发生时三本账全绿**——成绩、绩效、产出质量同时改善。认知负债研究不涉及这条制度性观察：它是实验室与调查设计，不是账本设计。

第二，**驱动变量不同**。卸载族的驱动变量是**使用强度**（用得越多越糟）；《底流》的驱动变量是**停喂时距**（多久没有亲手完成过一次带真实后果的执行）。两者可以背离：一个天天用 AI、却每天仍亲手完成若干带赌注执行的人，卸载族预测他退化，《底流》预测他底座保持。

判决形式：同一批人上并置测量使用强度、停喂时距与下游无辅助表现。若使用强度是唯一显著预测变量、停喂时距无独立贡献，则《底流》被卸载族吸收，第三部降级为该族的一个应用；若停喂时距独立预测且更强，则两说各管一段，且本书的账本处方（在能力清单上加“最近一次无辅助亲手执行”字段）有独立价值。

诚实还要求补一句对本书有利、却不能拿来当护身符的事实：认知负债那项研究本身已受到方法学质疑——有评论指出其部分脑电结果的一致性不足，并呼吁未来研究采用更严谨的设计、更大更多样的样本与更清晰的认知过程操作化。这条不能被本书用来贬低对手，只能用来提醒：这一族与本书都还没有跑完自己的账，第十六章第五场实验对两边同等重要。

四之八 品味与代理工程：2026 年的最新一族

主张。这一族出现得最晚、声音最响，且直接冲着“什么让人还值钱”来。品味一侧：Mollick 在 2026 年春的一次访谈中说，当网上多数写作都由 AI 生成、一切开始听起来一样时，能以独特风格写作的人反而更有意思；“品味是不是 AI 时代的关键差异点”已成为 2026 年科技界的一个争论话题。工艺一侧的表述更具体：Karpathy 在 2026 年 4 月提出“代理工程”，把它定义为一门编排易错的、随机的 AI 代理而非被动接受生成结果的职业学科，核心技能包括规格设计、逐差审查、评测设计、安全监督与质量品味；他并给出一条与本书第三部完全同向的判断——当 AI 变强，人的诱惑是学得更少，而他主张相反：理解成为瓶颈，你仍需足够的深度去指挥系统，知道该问什么、该检查什么、该拒绝什么、什么才要紧。

重合。重合处大到必须逐条点出：“逐差审查”就是《脱稿步》所说的、在产物延长线上的当场追问；“评测设计”就是本书通篇在做的事——自建判据并预注册；“理解成为瓶颈”与《底流》的存续条款方向完全一致。这一族是十二族里与本书距离最近的。

分界。分歧在**这是干活的学问，还是读人的学问**。代理工程与品味说的是**怎么把活干好**（生产工艺），本书说的是**第三方凭什么读得出谁把活干好了**（证据制度）。这个差别不是学究式的，它有两条可检验的后果：

其一，**品味一旦可命名，就进入批量供给区**。这是第二章第三通道对这一族的直接预测：当“风格独特”成为可教、可模仿、可提示词化的目标，它的差异度半衰期开始缩短（《零判份》的 ρ 随暴露单调收窄）。品味说目前没有关于“如何抵抗自身被批量供给”的机制，本书有一个：品味的读数只有在他择时刻、跨配组条件下才不可排演。

其二，**代理工程没有取样理论**。它精确回答了人该做什么（写规格、审差异、设评测），却没有回答一个雇主的问题：怎么知道**这个应聘者真的在做这些**，而不是准备了一份漂亮的代理工程作品集？规格与评测都是可预先准备的产物；一旦它们成为招聘信号，就会被排演——这正是《他手无事账》的整章内容。

判决形式：让一批人提交自选的代理工程作品集（规格、差异审查记录、评测设计），同时对同一批人做无预告的他择时刻取样。若作品集与他择读数对入职后表现的预测力相当，则本书的取证层多余，品味／代理工程一族完胜；若只有他择读数保留预测力，则这一族回答的是“该做什么”，本书回答的是“凭什么信”，两说互补而不互相吸收。

四之九 信号理论与文凭社会学：本书那句“产物不再是证据”的正主

主张。这一族比其余各族都老，而且它管的恰恰是本书的起点。信号理论说：教育的经济价值未必来自它教了什么，而来自它**发出的信号**——文凭之所以值钱，是因为它对能力强的人来说成本更低，因而能把人分开。分得开，靠的是“难以伪造”这个前提。文凭社会学那一支走得更远：文凭的功能主要是**筛选与封闭**，是群体争夺位置的凭据，而非能力的度量；于是会有文凭通胀——同一个岗位要求的学历一路上升，而岗位本身没有变。

重合。极大，而且是本书必须最先认下的一处。“**某类证据一旦变得容易伪造，就不再能分开人**”——这句话是**信号理论的自带推论，不是本书的发现**。本书第一章讲的“产物与能力之间的推断链断了”，在这一族的语言里就是一句现成的话：分离条件被破坏，信号失效。谁若说本书发现了这件事，是抬举了本书。

分界。分歧在失效之后往哪里去。

信号理论的处方沿着“提高伪造成本”走：找一个更难伪造的信号，或提高既有信号的门槛。文凭通胀就是这条路的自然结果——它是这一族自己预测的结局。本书主张这条路在今天走不通：**当准备可以被整体外包，任何“自择时刻交出的产物”都可以被伪造到同一水平，提高门槛只会提高所有人的成本而不改变分离度**。本书往另一个方向走：不去找更难伪造的产物，而是换到产物之外的四个维度——位次（答案何时到）、当场（是不是现配）、他择（时刻谁选）、记账（记在谁账上）。这四样都不是“东西”，因而不能被制造成东西。

判决形式：观察 2026 年之后各行业招聘的新增筛选环节。若新增的主要是**更难的自择产物**（更长的作品集、更难的笔试、更多的证书）且判别效度回升，则信号理论对，本书的**四维路线**多余；若新增的主要是**改变取样结构的环节**（无预告、现场续做、跨配组记录）而效度只在这些环节上回升，则本书对。这一条不需要实验，只需要三年的招聘公告与效度数据。

必须写下的一处不利。这一族里“文凭通胀”那一支对本书有一记重击：**凡是被制度采纳的分离手段，最终都会通胀**。若本书的四个维度真的进入制度，它们大概率也会通胀——出现“他择证明”的市场、“事件账”的代办、“现场续做”的培训班。本书没有反驳，只有第十七章那条元规则：谁采用本书的读数，必须同时采用该章的反噬条款。这不是解药，是限量的缓冲。

四之十 效度理论：本书最危险的一位占位者

主张。这一族来自心理与教育测量，它问的问题与本书的竖线几乎撞在一起：**一个分数的解释，在什么条件下才成立？**现代效度理论的核心主张是，效度不是测验本身的属性，而是**对分数所作解释与用法的属性**；判断效度要看支持该解释的证据与理论，还要看使用它所

带来的社会后果。后来的一支把它整理成"论证式"的做法：把从"作答"到"结论"的每一步推论都写成一条可以被质疑、可以被举证的链条，逐环检验。

重合——这里必须说得非常清楚。这是十二族里离本书最近的一族，比品味与代理工程那一族更近。本书第二章那条竖线主张"读数只在五处有定义"，用效度理论的语言几乎可以逐句翻译：那是对分数解释的**边界条件**的主张；本书的十个读数与各自的证伪条款，形式上就是一条条效度论证；本书第十七章关于"读数一旦进考核就被制造"的整节，在这一族里有现成的名字——**后果性效度**与指标被当作目标之后的败坏。

这一族在第四章的前八节里没有出现，是本书的一处真实疏漏，此处补记，并按纪律照登：它不是被本书驳倒的，是被本书漏掉的。

分界。认下这么多之后，剩下的差别只有一条，但它是实的：

效度理论问的是"**这个解释成立吗**"——它假定被测的那个属性存在，问的是我们的推论链有没有断。本书问的是更前面的一步：**在这些条件之外，那个量根本没有定义**。这不是效度不足，是定义域为空。用效度理论的框架，你会去检查"从单回合作品到个人能力"这条推论的每一环，然后发现它证据不足；用本书的框架，你会直接说：**单回合的账本里没有"个人"这个量**，无论你补多少证据，那一环都不会成立——就像除以零，不是精度问题。

第二条差别是内容上的：效度理论是一套通用的方法学，它不告诉你被测的属性长什么样。本书给出了属性侧的内容——发生的四条本体性质、一条存续条款、五处可读条件、十个读数。**方法学与内容不构成竞争**：本书的十个读数完全可以、而且应当被写成效度论证的形式来检验；如果有人这么做，本书受益。

判决形式：把本书的任一读数（例如他择率）按论证式效度的规程完整地做一遍。若走完全部环节之后，"单回合自择产物→个人能力"这条推论只是**证据不足**、可以靠补证据修复，则本书的"定义域为空"是夸张，本书应退回效度理论的框架内，作为它的一个应用；若无论补多少证据，该推论在结构上仍不成立（因为读数属于配组而非个人），则本书那条更强的主张站得住。

这一节对总账的影响，写在四之十四：本书原先自认的三处残量，有一处必须当场收窄。

四之十一 专长研究与刻意练习：与《底流》《现配优势》的正面相邻

主张。这一族问的是"高手是怎么长出来的"。一支强调**刻意练习**：专长来自长期、有明确目标、有即时反馈、在能力边缘反复进行的训练，而不是单纯的经验累积——干了二十年不等于练了二十年。另一支从真实现场入手，研究消防指挥、急诊、军事指挥这类高压环境中的**直觉性决策**：专家往往不比较方案，而是识别情境、直接生成一个可行动作；这种识别只有在**有反馈、有规律**的环境里才能长出来。

重合。与本书第三部（《底流》）几乎贴身。"干了二十年不等于练了二十年"和本书的"在喂率"说的是同一类事实；"识别只有在有反馈的环境里长出来"与本书的三成分判据（无辅助、真实后果、亲手）方向一致。与《现配优势》也近：现场识别本身就是当场配组的一种。

分界。两条，都可判负。

第一，**这一族的单位是训练，本书的单位是供给流**。刻意练习是一个人**特意安排**的活动——有目标、有反馈、有教练；本书的底流恰恰相反：它是日常工作里低价值部分的**联合副产**，从来不是任何人特意安排的，也从来没有被谁当作专款拨付过。这个差别有一个直接后果：按刻意练习的处方，能力衰退可以靠“多安排训练”补回来；按本书的读法，补不回来——因为断的不是训练，是嵌在真实赌注里的执行流，而训练不满足“真实后果”这一成分。

第二，**这一族关心能力怎么长，不关心能力怎么被第三方读出**。这与全书其余各节的分工一致：长与读是两件事。

判决形式：取一批底流已断（ β 很低）的在职者，分两组。甲组按刻意练习的处方补训（明确目标、即时反馈、能力边缘、有教练，但任务无真实后果）；乙组不补训，只恢复日常中带真实赌注的亲手执行。若甲组的下游无辅助表现恢复不逊于乙组，则《底流》的“真实后果”这一成分是多余的，第三部应退回刻意练习框架；若只有乙组恢复，则三成分判据成立，且这一族的处方在 AI 时代有一个盲点——它默认真实后果的执行机会一直在，而这正是被接走的东西。

四之十二 任务自动化经济学：为什么“还行的自动化”最伤人

主张。这一族用任务而非职业作为分析单位：技术既有**替代效应**（把任务从人手里拿走），也有**恢复效应**（创造新任务把人放回去）；两者的对比决定劳动份额与工资走向。它还给出一个特别有解释力的概念——“**还行的自动化**”：那些只是勉强比人做得好、生产率提升有限的自动化，最伤人。因为它足够便宜到把人替换掉，却没有好到能把蛋糕做大、创造出足够的新任务来把人接住。

重合。它与《峰兑》的地板抬平、《底流》的代产断供在机制上同源，而且它提供了本书缺的那一半——**宏观动因**。本书解释的是个人读数为什么塌，这一族解释的是为什么会有那么多“刚好够便宜、刚好够用”的替代发生。

分界。单位不同，这一点与图灵陷阱那一节相同：它的单位是任务与劳动份额（宏观），本书的单位是读数的定义域（微观、证据）。但本书在这里有一个可以还回去的东西，而且这是本章少见的一处**正向贡献**：

这一族衡量替代的尺度是**生产率**——自动化好不好，看它把产出提高了多少。本书的第九章提出了另一把尺子：**同一次替代，即使生产率提升很大，也可能同时切断底流**。换句话说，“还行的自动化”之所以伤人，除了它没做大蛋糕之外，还因为它接走的往往正是那些低价值、重复、无人记账、却恰好在喂人的执行。**按生产率算，这类替代常常是划算的；按在喂率算，它可能是最贵的。**

判决形式：在同一批岗位上并置两把尺子——替代带来的生产率变化，与被替代者的停喂时距变化。若两者高度相关（生产率提升越大、停喂越少），则本书这把尺子是多余的，被这一族吸收；若两者可以背离——存在生产率大幅提升而停喂时距同时恶化的岗位——则底流提供了一个这一族尚未纳入的成本项，而这个成本项目目前不出现在任何一份自动化投资的测算里。

四之十三 来自诸说的四记反诘，与本书的答复

划界是本书对外说话；这一节反过来，把十二族对本书最锋利的四记反诘集中列出，并把答复照实写下——其中两记，本书答不上来。

反诘一（来自文凭社会学）：凡是被制度采纳的分离手段，最终都会通胀。你的四个维度若真被采用，就会出现“他择证明”的市场、事件账的代办、现场续做的应试培训班。

答复：本书没有反驳。这一记打中的是所有评价制度的共同命运，本书不例外。能给的只有两样：其一，第十七章的元规则——采用任一读数者必须同时采用该章的反噬条款，二者不得拆开引用；其二，把所有权挪到人这边——记录归被读者本人持有、不进排名与分配，则伪造它只是骗自己。这两样都不是解药，是缓冲。这条反诘作为本书的一个开放弱点登记在案。

反诘二（来自效度理论）：你所谓“定义域为空”，多半只是效度证据不足的一种夸张说法。

答复：本书接受它是一个真正的、可能致命的质疑，并在四之十写下了判定形式：按论证式效度把任一读数完整走一遍，看那条推论是“证据不足可修”还是“结构上不成立”。这一场未跑。在跑完之前，本书那条更强的主张只能算作待检的猜想，而不是已立的结论——第二章的措辞因此在再版时可能需要收窄。

反诘三（来自品味与代理工程一族）：你的他择时刻同样会被表演——说好无预告，实际上内部提前通气。

答复：会。这是全书最难防的一处，本书防不住。能做的只是把它从一个不可见的问题变成一个可查的舞弊：编码规则公开、取样时刻由不参与本次评价的第三方决定、留申诉通道。这些手段都不新鲜；要紧的是先把“取样时刻由谁定”当作一个对象，才会有人去设这些防线。现在的困境不是防不住，是没有人防。

反诘四（来自专长扩展派与任务自动化经济学的合流）：如果专家任务本身在萎缩，你那五处可读条件的定义域也在收缩。跨配组要有第二个班子，他择读数要有人愿意读你，事件处要有事件落到你头上——这些全是机会，不是能力。

答复：这一记本书接得住，但接得不体面。本书承认：发生差首先记录的是曝露的分布，其次才与能力相关。一个从未被读过的人，读数为零，价值不为零。这就是第十七章第一条伦理底线的由来，也是本书唯一一处把“机构必须先自证”写成硬条款的地方：任何机构以本书读数作不利安排之前，必须先回答“此人有没有被读的机会”，答不出，安排无效。但这条底线管不住机会本身的分布——那不是一本书能解决的问题，本书也不假装能。

四之十四 总账：本书被占位到什么程度

把十二族并排放好，本书的位置才可以如实标定。这一节在本书付梓前改写过一次——四之十那一族（效度理论）是补记的，它一进来，本书原先自认的三处残量当场少了半处。改写的痕迹保留在这里，不抹平。

残量一（已收窄）：可读条件是定义域声明，不是测量建议。十二族里，有十一族默认差异随时可测、只是难测。唯一的例外是效度理论：它明确主张分数解释依赖条件，并要求逐环举证。因此本书这一处残量必须从“无人占位”收窄为“与效度理论共占，且本书主张更强”——本书说的是“条件之外那个量没有定义”，效度理论说的是“条件之外那个解释缺乏证

据"。这两句话不同，但差别有待四之十那场检验来裁；在它跑完之前，本书对这一处的所有权只能算**待定**。

残量二：取样时刻的归属。十二族里无一把"这份证据的取样时刻由谁决定"当作对象。信号理论关心信号难不难伪造，却不问时刻谁定；效度理论关心推论链，却把取样条件当作给定；专长研究关心训练环境，同样不问。这是《他手无事账》的整章命题，也是本书唯一一处**发现空格**的地方——三个行业各花了一代人，八格里只填到七格。

残量三：账本载体的二分与记账权。十二族都在谈能力，无一谈能力记在谁的账上；文件载体与名单载体的归属规则不同，且 AI 只能吞下前者——这一条来自一百多年的判例，不来自任何关于人工智能的讨论。

除这两处半之外，本书必须承认的占位相当密，逐条点名：**"产物一旦易于伪造即失去分离力"属于信号理论**；地板抬平被参差前沿一族先测；断供机制被认知负债一族先测（虽然驱动变量不同）；"干了二十年不等于练了二十年"属于专长研究；配组归属在读数分解与团队生产传统里早有位置；"廉价者被接走、互补者升值"的骨架属于预测—判断二分；宏观动因分属图灵陷阱与任务自动化经济学；而关于**"分数解释在什么条件下成立"这一整套问法，属于效度理论**。

本书在这十二族里，**不是一个新发现，是一次重新对焦**：把十一族各自默认"随时可读"的那个量，摁在它真正有定义五个位置上，给出十个读数与各自的证伪条款，并补上两处此前无人当作对象的东西——**取样时刻的归属，与账本载体的记账权**。

最后是本章自己的三条对赌，写死不改：

最后是本章自己的三条对赌，写死不改：

其一（吸收之赌）：若上述任一族在其后续工作中独立提出"取样时刻归属"或"账本载体二分"并给出读数，本书的残量二或残量三即被占位，第一章"缺口是踩出来的"随之降级——不是被推翻，是被证明该缺口早有人从别的方向走到，本书只是晚到者。

其二（定义域之赌）：若在 2031 年 12 月 31 日之前，出现一种在**单回合、平时、自择时刻、已入账**条件下稳定预测入职后表现的评价方法（效度达到 2022 年前水平并被独立复现），则本书的可读条件全部为伪——不是不够精确，是从根上判错了定义域。这条与第三章失效条件三互为表里，一并作废。

其三（复检之赌）：本章每十二个月重扫一次。任何一条分界若被对手方自己的新工作抹平，即时公告作废；本书不为已被抹平的分界辩护。**一部主张"发生差要靠被检验的执行来喂养"的书，它自己与外界的分界也必须被反复重走，否则这一章就是一份自择证据——写作时刻由作者选，材料由作者挑，正是本书通篇警告的那种东西。**

本章所涉主要文献

（按出现次序；网络来源的访问日期均为 2026 年 9 月 3 日。）

- World Economic Forum, *Future of Jobs Report 2025*, Geneva, January 2025.
- 曾湘泉等，《人工智能时代，我们如何更好就业创业》，《人民日报》理论版，2026 年 4 月 9 日。
- 《全球数字教育发展指数（2026）》，2026 世界数字教育大会发布；相关报道见《人民日报》2026 年 5 月 31 日第 5 版。

- Autor, D. H., "Applying AI to Rebuild Middle Class Jobs," NBER Working Paper No. 32140, February 2024 (另以《AI Could Actually Help Rebuild The Middle Class》刊于 *Noema*, 2024 年 2 月 12 日)。
- 英国等地劳动力市场的相反机制证据, 见 "Generative AI and Labor Market Outcomes: Evidence from the United Kingdom" 一族工作论文及其对 Autor (2024) 的评述。
- Agrawal, A., Gans, J., & Goldfarb, A., *Prediction Machines: The Simple Economics of Artificial Intelligence*, Harvard Business Review Press, 2018 ; 同著者 *Power and Prediction: The Disruptive Economics of Artificial Intelligence*, Harvard Business Review Press, 2022 ; 并见 "Exploring the impact of artificial intelligence: Prediction versus judgment," *Information Economics and Policy* 47 (2019): 1-6.
- Brynjolfsson, E., "The Turing Trap: The Promise & Peril of Human-Like Artificial Intelligence," *Dædalus* 151(2), Spring 2022, pp. 272-287.
- Dell'Acqua, F., McFowland III, E., Mollick, E., Lifshitz-Assaf, H., Kellogg, K., Rajendran, S., Kraymer, L., Candelon, F., & Lakhani, K. R., "Navigating the Jagged Technological Frontier: Field Experimental Evidence of the Effects of Artificial Intelligence on Knowledge Worker Productivity and Quality," Harvard Business School Working Paper 24-013, September 2023 ; 正式发表于 *Organization Science*, March 2026.
- Kosmyrna, N., Hauptmann, E., Yuan, Y. T., Situ, J., Liao, X.-H., Beresnitzky, A. V., Braunstein, I., & Maes, P., "Your Brain on ChatGPT: Accumulation of Cognitive Debt when Using an AI Assistant for Essay Writing Task," arXiv:2506.08872, June 2025 (v2, December 2025)。
- Stanković, M., Hirche, E., Kollatzsch, S., & Doetsch, J. N., 对上文的评论 (arXiv, 2026)。
- Gerlich, M., "AI Tools in Society: Impacts on Cognitive Offloading and the Future of Critical Thinking," *Societies* 15(1), 2025.
- Mollick, E., 关于"品味"的访谈言论, 2026 年 4 月 30 日播出; 相关报道见 2026 年 5 月。
- Karpathy, A., "Agentic Engineering", Sequoia Ascent 2026 演讲及作者本人的讲稿摘要, 2026 年 4 月 30 日。
- Spence, M., "Job Market Signaling," *Quarterly Journal of Economics* 87(3), 1973: 355-374.
- Collins, R., *The Credential Society: An Historical Sociology of Education and Stratification*, Academic Press, 1979.
- Messick, S., "Validity," in R. L. Linn (ed.), *Educational Measurement*, 3rd ed., Macmillan, 1989, pp. 13-103.
- Kane, M. T., "Validating the Interpretations and Uses of Test Scores," *Journal of Educational Measurement* 50(1), 2013: 1-73.

- Ericsson, K. A., Krampe, R. T., & Tesch-Römer, C., "The Role of Deliberate Practice in the Acquisition of Expert Performance," **Psychological Review** 100(3), 1993: 363–406.
- Klein, G., **Sources of Power: How People Make Decisions**, MIT Press, 1998.
- Acemoglu, D., & Restrepo, P., "The Race between Man and Machine: Implications of Technology for Growth, Factor Shares, and Employment," **American Economic Review** 108(6), 2018: 1488–1542 ; 及同著者 "Automation and New Tasks: How Technology Displaces and Reinstates Labor," **Journal of Economic Perspectives** 33(2), 2019: 3–30 (含"还行的自动化"一说)。

方程的右端从"发生"二字开工，因为其余一切——供给、读数、取证、归属——都是围着这两个字的后勤。本部四章按一件事的时间轴排开：**答案到达之前**（《盲遍》——那一遍要么在答案之前完成，要么没有完成）、**第一次**（《单遍段》——有些工序走过一次就耗散了重来的条件）、**交手当场**（《现配优势》——显露是当场配组做出来的，不是从存量调取的）、**交手之后**（《脱稿步》——发生留下的署名，须在自己产物的延长线上一次次续签）。

四章合起来给"发生"画出的轮廓，第二章二之二已用一句话收拢：答案之前、恰有一遍、当场配成、须被续签。读这四章时值得带着的一个对照是：这四条没有一条依赖"人工智能"这个词成立——谜语在语言模型出现之前就不能猜第二次，舒芙蕾在自动化厨房出现之前就塌了回不去，评酒与育种的复样数据躺了几十年。四章各自的第一节都会交代，AI 改变的只是这四条从"读不读得出无所谓"变成"唯一还立着的东西"。

四章各自的欠账也如实预告：四场旗舰检验（随机化答案位次、跨域重做样本库、交接复现实验、判者一致性预实验）至本书付梓无一真跑，协议全部备妥，收在第十六章。

第二部 发生：本体层

第五章 盲遍

【章首】本章给方程添的那一刀，是发生的**位次性**：对每一个具体未知，每人恰有一遍“答案不在场时的成形承诺”——答案先到，这一遍即被无声取消，档案上一字不差，事情却没有发生过。读数是门位次 k ：答案到达落在承诺序列的第几步。它与《脱稿步》最近，分界已写进两章正文：脱稿步读事后被发起的追问，盲遍读形成期的答案时序——档案与追问全同的两人，形成期位次分布不同则后续表现不同，是两章的判决实验。本章的证据包括两场事先登记的检验（盲法元研究统计量的再判读、开放徽章 3.0 全部一百六十三个字段的审计），以及一处已在正文坦白的占位：后见之明研究的标准处方“结果揭晓前写下预测”在操作层几乎覆盖本章的旗舰形态——三条分界与可分实验随坦白同页。本章欠的账：随机化答案位次的真实迁移实验（协议在正文，任何一门课可执行），未跑。

大白话：有些事，你只有一次机会自己想

先说一个所有人都懂的例子：猜谜语。

有人给你出一个谜面，你想了半分钟，说出一个答案。这半分钟里发生的事，叫“猜”。现在换个情形：有人先告诉你谜底，再把谜面念给你听，问你“你猜得出来吗”。你当然“猜得出来”——但你我都知，那不是猜。你并不是猜得不好，是“猜”这个动作，对这个谜语来说，对你已经不存在了。

这就是《盲遍》要说的全部事情。它给这件事起了个名字：对每一个具体的未知，每个人恰好只有一遍“答案不在场时的成形承诺”。这一遍用掉了，就没有了；答案先到，这一遍就被无声地取消——注意“无声”两个字：档案上一模一样，卷面一模一样，别人看不出，你自己也很容易忘记，但那件事确实没有发生过。

为什么这在今天忽然变得要紧？

因为在 2023 年之前，“答案先到”是一件需要成本的事。要抄，得有人抄；要查，得查得到；很多问题根本没有现成答案。于是绝大多数人在绝大多数问题上，天然地拥有那一遍。它太廉价了，廉价到没有人觉得它是资产。

现在，任何一个具体问题的解析，都可以在三秒内到达你眼前。于是那一遍从“天然拥有”变成了“必须刻意保护”。你人生中还剩下多少个具体未知，是你打算自己先走一遍的？这个数在飞快地掉，而且掉的时候没有任何声音——每一次“我先问问它”，都是从这个账户里划走一笔，而账户上不打印流水。

它不是什么，这点必须说清楚。

它不是“不许用工具”。用不用工具，是效率问题；这一遍在不在，是那件事有没有在你身上发生过的问题。一个人可以每天用十个小时的工具，同时守住自己那一遍；也可以从不用工具，却因为总是先翻答案，把那一遍全花光。

它也不是“记忆力”或者“独立思考”这类旧词。旧词说的是能力的强弱，这里说的是**位次**——同一件事，在答案到达之前做，和在答案到达之后做，是两个不同的动作，不是同一个动作的两个水平。

读数是什么？

书里给它一个可以数出来的量，叫**门位次 k** ：把你面对一个具体未知时的动作排成序列——想、试、写下判断、查、问人、看答案——答案到达落在第几步， k 就是几。答案第一

步就到， k 等于零，这一遍没有发生；你先形成了一个成形的、可以被自己写下来的判断，然后才去查， k 就大于零。

请注意“成形”两个字：脑子里模模糊糊转过一下不算，那个念头必须被**写下来或者说出口**，成为一个可以被后来的事实反驳的东西。这是全书反复出现的一条纪律：**一个不能被反驳的判断，等于没有判断过。**

一个小场景。

两个新入职的年轻人，都要写一份市场分析。甲拿到题目，先花二十分钟，把自己认为的三条关键因素写在纸上，标了自己最不确定的一条，然后才打开工具，用两个小时把报告做完。乙拿到题目，直接让工具出一版，然后花两个小时把它改得很漂亮。

两份报告交上来，可能乙的那份还更好看。档案上看不出任何差别。但半年之后，公司遇到一个从来没有人写过分析的新问题——那时候甲手里有二十分钟里练过的东西，乙没有。这不是甲更聪明，是甲那一遍还在，乙那一遍已经被换成了一份成品。

最后说一句它的边界。

守住那一遍，代价是慢，而且经常错。所以它不该被用在所有事情上——很多事情就直接查答案，人生没那么多必须自己发明轮子的时刻。它值得被守住的地方是：**那些你打算靠它吃饭的领域里的具体未知**。在别处慷慨，在这里吝啬。

再看两个身边的例子。

学开车。教练在旁边，每个路口他提前半秒说“打方向”，你照做，一路很顺。你觉得自己会开了。等到某天单独上路，遇到一个教练没讲过的路口，你才发现自己练的是“跟着提示做动作”，不是“自己判断”。教练那半秒，就是提前到达的答案。他不是坏教练——恰恰因为他是好教练，他让你开得很顺，而顺本身掩盖了那一遍没有发生。

看侦探小说。同一本书，第一次读和知道凶手之后再读，是两本不同的书。第二次你能看出所有伏笔，你的“理解”甚至更深——但那个“边读边猜、猜错、修正”的过程，永远不会再有第二次了。这就是为什么剧透让人生气：被拿走的不是信息，是那一遍。

三种常见的误读。

第一种：把它当成“要靠自己”。不是。这一章从不反对求助、查资料、用工具，它只管一件事——那些动作发生在你形成判断之前还是之后。先形成判断再去查，查得越多越好；先看答案再来“想一遍”，想得再久也不算。

第二种：把它当成学习方法。它比学习方法更基础。学习方法讲的是怎么学得更好，这一章讲的是**某个特定的动作有没有在你身上发生过**。你可以用最糟糕的方法守住那一遍，也可以用最先进的方法把它全部花光。

第三种：以为可以补。补不了。这是它最要紧、也最容易被忽略的性质。技能可以补，知识可以补，那一遍不能——因为它的条件是“答案不在场”，而答案一旦到场就不会再离场。这就是为什么它值得被当作资产来管理：**世界上大部分资产花掉了还能再挣，这一个花掉就没有了。**

给自己做个小检查。回想你最近一个月里认真处理过的五件事，逐件问：动手之前，我有没有写下过一句可以被后来的事实推翻的判断？五件里有几件？这个数就是你现在的样子。

它和旁边几个概念的区别。

和《单遍段》最近，但两者管的是两扇不同的门：这一章管的是**你和答案的关系**（答案不能先到），《单遍段》管的是**你和对象的关系**（过程不能重来）。一道纯判断题，答案先到就触发这一章，而对象一点没动，与《单遍段》无关；一次不可撤销的操作，哪怕世界上根本没有标准答案，也照样触发《单遍段》。

和《脱稿步》的分工也清楚：《脱稿步》读的是事后被人追问时你还接不接得上，这一章读的是形成期答案是第几步到的。两个人可以档案一样、追问表现一样，而形成期的位次分布完全不同——那正是两章的判决实验。

和《底流》的关系最容易混：底流问“你最近有没有亲手做”，这一章问“你做的时候答案在不在场”。**一个天天亲手做、但每次都先看答案的人， β 很高而 k 恒为零。**这两个量必须分开看，只看一个都会得出错的结论。

一句话记住它：这一章管的是次序——你的判断和答案，谁先到。

怎么长出来：方法、技术与流程

个人层面（三件小事，成本几乎为零）。

第一，**承诺先行**：在自己的专业领域里，凡遇到一个具体未知，先写下一句可被反驳的判断，再动手查。一句就够，写在便签、聊天窗口给自己发一条都行。关键不是写得好，是**先于答案存在**。

第二，**留位次记录**：每周挑三件事，简单记一下 k 是几（答案第几步到的）。一个月之后你会看到自己的分布——多数人会把自己吓一跳。

第三，**分区**：明确划出“这块我要自己先走一遍”和“这块我直接要答案”。不划分区的人，最后会在所有事情上都直接要答案。

大学发生学教育：把“那一遍”做成课程的一部分。

这一条的设计要点是：**不禁止工具，而是安排位次。**

- **承诺卡（每门课每学期十次，五分钟）**。发下题目后，学生先用五分钟在一张卡上写下自己的初判与最不确定处，卡交上或拍照上传，**然后**教师宣布“工具开放”。承诺卡不按对错评分，只按“是否成形、是否可被反驳”评分——写“我认为可能有很多因素”得零分，写“我认为主因是 A，如果 B 的数据不支持，我就错了”得满分。
- **答案延迟发布制**。参考答案、标准解、范例作业，一律在承诺环节完成之后才发布。这不需要新系统，只需要教务把“发布时间”当作一个教学变量来管。
- **迁移考核代替复现考核**。期末不考已经练过的题型，考一个结构相同但从未出现过的新问题。守住过那一遍的人，和把每一遍都换成成品的人，在这里第一次分开。
- **位次字段进档案**。学生的课程记录里加一栏：本学期承诺先行的次数。它不参与评优，也不进排名——**它一旦进排名，第二天就会有人批量刷承诺卡**（这是全书反复的警告：读数一旦成为考核目标，它读的那件事就开始被制造）。它的用途只有一个：让学生自己看见自己的账户余额。

教师的三个具体动作。

其一，改口头禅：把“这道题答案是什么”换成“你先说说你打算怎么走”。其二，公开自己的位次：教师在课堂上遇到不会的问题时，当场演示“我先猜一个，再去查”，而不是回避或直接投屏搜索。其三，**保护错的承诺**：承诺错了不扣分，只记录；一旦承诺错要扣分，学生就会等到有把握再写，那就等于答案先到。

最容易做坏的地方。

把承诺卡变成打卡任务，是这条设计唯一的死法。防它的办法只有一个：**永远不把承诺的次数与任何利益挂钩**，只把它当作学生自己的账本。

原文

盲遍——人工智能时代大学毕业生的核心竞争力，是答案到场之前完成的那一遍

副题的可裁决形式：对每一个具体未知，每人恰有一遍"答案不在场时的成形承诺"；答案先到，这一遍即被取消，且取消不留任何档案痕迹。据此，毕业生之间可持久的竞争差，读在答案到达的位次上，而不读在档案的内容上。

摘要

"人工智能时代，大学毕业生的核心竞争力将是什么"——现成回答不外某张技能清单、某种信号或某种学习能力。本文指出三类回答共享一条未被说出的地基：一个人此刻能显露什么，由他此刻拥有什么决定；多拥有一样东西（信息、经历、答案），至少不会减少他能作出的显露种类。来自盲法研究自身的实证推翻了这条地基：知情判读者拥有更多信息，却永久失去了产出盲读数的能力，且两态判读之差随结局主观度与判读者卷入而变化——知情后的判读不是同一动作加上更多信息，而是另一种动作。承重命题：每人对每一具体未知恰有一遍"未知在场时的成形承诺"，称为盲遍；答案先到即预走，预走无声、不损档案、不可赎回。判据：取一件交付，核对答案可得时刻与承诺落定时刻的先后；答案后到，盲遍成立；答案先到，该件交付上不存在判断的形成。旗舰读数为门位次，即答案到达落在承诺序列的第几步。两次事先登记的检验均未触发判负条款：盲法元研究统计量的再判读使"另一种动作"层获得弱支持；开放徽章 3.0 规范一百六十三个字段的审计显示，七个时间字段无一由持有者触发。结论一句：人工智能时代大学毕业生的核心竞争力，是他名下仍在发生的盲遍，以及把答案排到承诺之后的位次安排。

关键词：盲遍；门位次；预走；全对而未走过；观察者偏倚；春化重置；后见之明偏差

Blind Pass: The Core Competitiveness of College Graduates in the AI Era Lies in What Is Committed Before the Answer Arrives

Abstract: For every specific unknown, each person gets exactly one pass at forming a committed judgment while the answer is absent. When the answer arrives first, that pass is silently cancelled, leaving records indistinguishable from genuine work. Against Plato's Meno paradox and its anamnesis solution, against productive-failure research (Kapur), the assistance dilemma (Koedinger and Aleven), and recent "protected first attempt" frameworks for AI-era pedagogy, this paper argues the issue is not the efficacy of struggle-first ordering but the existence of a one-way passage: evidence from blinded-versus-nonblinded outcome assessment (Hrobjartsson et al.) shows the informed reading is a different act, not a corrigible offset, echoing the single-use structure of Shannon's one-time pad. The flagship measure is gate position: at which step of the commitment sequence the answer arrives. Three formal definitions and three propositions (uniqueness, non-redeemability, positionalization) are given, together with a round-based sequence model, a fourfold typology of transcripts whose target cell is 'all correct, never walked', two pre-registered tests that

survive their falsification clauses (a re-reading of published meta-research statistics, and a field-by-field audit of the Open Badges 3.0 schema showing none of its seven temporal fields is holder-triggered), a full executable protocol for the decisive randomized-position experiment, an explicit demarcation from the hindsight-bias tradition (Fischhoff) whose standard remedy—write predictions before outcomes—coincides operationally with this paper's prescription while differing in what the act is for, and mechanism testimony from vernalization resetting (Sheldon et al.; Crevillen et al.) and antigenic imprinting (Francis; Gostic et al.). Implications are drawn for human capital accounting, for test theory (formation accounting), and for the temporal design of human-AI division of labor.

Keywords: blind pass; gate position; pre-walking; observer bias; vernalization reset; original antigenic sin; hindsight bias

一、引言

设想一位面试官面前放着两份毕业生档案。成绩单相同，课程项目相同，代码仓库的提交记录一样干净，作品集里连排版都难分高下。按现行的一切读法——成绩、证书、作品、测评——这两个人无法区分。

再设想这两份档案的形成过程有一处不同。第一位毕业生做每道题、写每段程序时，先在没有任何解答在场的条件下把自己的判断交出去：写下答案、提交代码、按下运行，然后才对照解析、才看模型给的版本。第二位毕业生的每一步都有解答先行：题目旁边就是解析，编辑器里自动补全先于他的构思出现，他核对、理解、确认，然后落笔。两人当期的正确率几乎一样——第二位甚至略高，因为解答先到确实减少了当期错误。

四年之后，这两份档案逐字相同。本文要说的是：这两个人不同，而且这个不同在档案上永远查不出来。问题只在于，那个不同究竟是什么、坐落在哪一步、用什么读出来。

这就是本文的问题与来路：问题从一切档案读法的失明处来；此刻问它，因为答案供给的默认位次正在被人工智能整体前移，失明第一次从个案变成产线。本文立的东西有三件：一个存在物（盲遍），一个读数（门位次），一条判据（承诺与答案的先后）。在方法上，本文做三件性质不同的工作，各有各的证据标准：其一是概念工作——给出盲遍、预走、门位次的显式定义与判据，其合格线是可机械执行、无情态词；其二是机制迁移——从盲法方法学、春化表观遗传学与免疫印记研究三处成熟文献中取回同构机制，其合格线是逐条注明出处与不同构之处；其三是一次事先登记的再判读——对已发表元研究统计量按预先写死的判负条款作再判定，其合格线是登记先于读取、不利结果照登。本文不含新的一手实验，位次层的直接实验检验以可执行方案的形式交在第十六节。路线图：第二、三节摆出日常现象与三条研究传统的共同预设，第四节以盲法研究自身的实证构成反例，第五至九节立概念、命题、测量与模型，第十、十一节从春化与免疫取回再生机制与锚定推论，第十二、十三节完成理论定位与同批分界，第十四至十六节交出登记检验、边界条件与可证伪预测，第十七至十九节给出实践含义、理论后果与伦理件，第二十节以局限、自反与一条写死日期的预测收束。

二、日常现象：一个谜语只能猜一次

先从一件人人经历过的小事说起。一个谜语，你一生只能猜它一次。听过答案之后，你还能“答对”它一万次，但你永远不能再猜它——不是变难了，是“猜”这个动作对你不再存

在。同理：被剧透的电影不能再第一次看；对过答案再做对的题，不叫会做；朋友把脑筋急转弯的谜底先说了，这道题就从你的世界里被划掉了，划掉它的不是恶意，常常是好意。

这件小事有三个容易被放过的特征。第一，它无声。谜底先到的那一刻没有任何警报，当事人的体验往往是“哦，原来如此”，甚至是“学到了”。第二，它不损档案。你答出谜底的记录，与你猜出谜底的记录，写在纸上一模一样。第三，它不可赎回。世界上没有任何价格能买回你的不知道——能用钱买回来的叫成本，买不回来的才叫单向门。你可以雇一个没听过谜底的人替你猜，但那是换人，不是赎回。

日常语言里早有对这件事的直觉：老人说“这层窗户纸不能替他捅破”，师傅说“先让他自己焊一遍再讲”。但直觉停在劝告上，没有变成判据。本文要做的是把它做成判据：坐标定在哪一步，读数是什么，什么证据能推翻它。顺带记下这条直觉在行业里的两处残留，后文要用：外科培训至今保留“先让住院医说出他的方案、再告知主刀的方案”的问答次序；围棋道场至今要求学生先摆出自己的应手、再看定式——两处都在守同一步先后，两处都说不出为什么非守不可。

三、文献综述：三条研究传统与它们共同的预设

对毕业生竞争力的既有研究，可归入三条各自成熟的传统。

第一条是人力资本传统。自贝克尔（Becker，1964）把教育刻画为对生产性存量的投资起，这条传统的当代形态就是技能清单：竞争力等于个体持有的技能组合，清单随技术变迁滚动更新。戴明（Deming，2017）对社会技能溢价上升的测算、世界经济论坛逐年发布的未来就业技能榜（World Economic Forum，2025），都是这条传统的延伸——榜单每年换行，“竞争力在于持有清单上的行”这个框架五十年未换。

第二条是信号与筛选传统。斯宾塞（Spence，1973）与阿罗（Arrow，1973）指出：教育的劳动市场价值可以与其生产性内容脱钩——文凭之所以值钱，是因为取得它的成本与能力负相关，雇主借以筛选。这条传统在人工智能时代的直接推论是信号贬值论：当模型能替任何人产出合格的作品，一切可展示物的筛选力同步下降，于是竞争力研究变成寻找“下一个还没被稀释的信号”。

第三条是任务与互补传统。奥托（Autor，2015）与阿西莫格鲁、雷斯特雷波（Acemoglu & Restrepo，2019）把职业拆成任务束，问哪些任务被机器替代、哪些与机器互补；毕业生竞争力被改写为“任务组合向互补侧迁移的能力”。这条传统近年拿到了两项被广泛引用的一手证据：诺伊与张（Noy & Zhang，2023）发现生成式模型使写作任务的产出差距被压缩、低水平者获益最大；布林约尔松等（Brynjolfsson，Li & Raymond，2023）在客服现场发现同样格局——新手生产率提升最多，老手几乎不动。两项研究的作者都按普惠读这组发现。本文将在第九节给出另一种读法：新手获益最大，恰是答案先到的当期签名——被压缩的不只是产出差距，还有新手判断力的形成本身；这两种读法对同一批数据作相反的五年期预测。

三条传统在校园与雇主的日常语言里各有通俗化身：清单派开列批判性思维、提示词工程、跨学科协作与审美共情；信号派紧盯学历与证书在甄别中的残值；元能力派宣称具体技能皆会过时，唯“学习如何学习”不朽。三派彼此吵得很凶：清单派嫌信号派空，信号派嫌清单派天真，元能力派嫌前两派都盯着存量。但三派脚下踩着同一条没人说出口的地基——

一个人此刻能显露出什么，由他此刻拥有什么决定；多拥有一样东西（一条信息、一段经历、一个答案），至少不会减少他能作出的显露种类。

清单派数拥有的技能，信号派读拥有的凭证，元能力派押拥有的获取能力——三派全部默认：占有对显露单调不减，多知道点总没坏处。正因为如此，三派对“答案的供给成本趋零”这件事都只能给出乐观读法：答案便宜了，等于人人多拥有了，等于起点普遍抬高了，竞争差只是挪到清单上更高的行去了。

这三派在人工智能时代各有当代变体，且都在走红。清单派的最新版本是“人机协作技能”：提示词写作、结果核验、 workflow 编排——把与模型打交道本身列成新的技能行。信号派的最新版本是“作品集通胀论”：既然文凭信号被稀释，就改押可展示的真实项目。元能力派的最新版本是“学习敏捷度”：模型三个月一换代，唯一不贬值的是换乘能力。三个变体各有道理，但请注意它们与母版共享同一个动作——都在“拥有”这一列里加行，没有一个追问过：答案的普遍先到，会不会取消某种拥有之外的东西。本文的全部工作就压在这一问上。

这条地基如此自然，以至于推翻它的证据长期躺在一门学科自己的档案里，没人把它当回事。

四、反例：知情判读是另一种动作

临床试验方法学里有一个成熟的分支，专门研究“判读的人知道了不该知道的事，会发生什么”。它的标准装置是盲法：把分组信息挡在结局判读者之外，直到判读落定。这门分支自一九四八年医学研究理事会的链霉素试验（Medical Research Council, 1948）起步，积累了自己的奠基文献、教科书与反例传统。

它最重要的一批实证，来自一种设计巧妙的比较：同一批试验、同一个结局，同时安排知情判读者与盲态判读者各判一遍。赫罗别雅特松等人（Hróbjartsson et al., 2012）对二十一类此类试验的系统评价给出三条读数。第一，知情判读者平均把疗效的比值比夸大百分之三十六（合并比值比之零点六四，区间零点四三至零点九六）。第二，同一份病人材料，知情读与盲读在约五分之一的个案上给出不同判定——一致率中位数百分之七十八。第三，也是最要紧的一条：夸大的幅度与结局的主观程度评分无关，与判读者在试验中的卷入程度无关，与结局对期望的易感性无关。

第三条读数值得停一停。如果知情只是“多了一条信息、于是判读被这条信息带偏”，那么带偏的幅度应当随结局的软硬、随判读者的利害而系统变化——软结局偏得多，硬结局偏得少，利害深者偏得多。实测不是这样：偏移不认这些最顺理成章的旋钮。二〇二五年的扩大更新（Salazar et al., 2025；六十六项试验）在更大样本上重复了这个格局：主观度、易感性、卷入度三条元回归全部空手，而药物试验与非药物试验的偏移方向竟然相反。

把这三条合起来读，结论只有一个形状：**知情之后的判读不是原来那个动作加上更多信息，而是另一种动作**。知情判读者拥有更多，却永久失去了产出盲读数的能力——整个盲法制度存在的前提，就是承认这种失去无法由本人凭意愿撤销：一个看见过分组的判读者，无论多么自律，都不能把自己重新变盲，只能被换掉。多拥有一样东西，让一种显露从此对他关门。

这就推翻了上一节那条地基，而且是用地基受益最深的一门学科自己的数据推翻的。占有对显露不是单调不减的：有一类占有——答案的占有——恰恰取消一类显露。

有人会问：这难道不能靠训练与自律解决吗？让判读者知道自己有偏，努力矫正就是。这条路方法学界走过，结论是走不通，理由有两层。经验层：舒尔茨等人（Schulz et al., 1995）对二百五十项试验的分析显示，分组隐匿不充分的试验平均把疗效夸大三到四成——而这些试验的判读者没有一个不知道偏倚的存在，知道不等于能免。机制层：矫正预设了一个稳定的偏移量可供反向抵消，而上文第三条读数说的正是不存在这样的稳定量——偏移不认主观度、不认卷入度、方向还会在亚组间反转。你无法校正一个不是偏移量的东西，正如

你无法把一次"答对"折算回一次"猜出"。方法学的最终解决方案从来只有一个：换人——用一个还不知道的人。这个"只能换人、不能洗净"，就是本文全部论证的支点。

需要与一个更著名的邻居分开：期望效应研究。罗森塔尔与雅各布森（Rosenthal & Jacobson, 1968）证明，教师对学生的预先信念会改变教师的后续行为并部分自我实现。期望效应说的是"知道改变了知道者对对象的作为"，作为可以被训练、被监督、被抵消——它是行为层的事。本文取自盲法实证的是更深一层：知道取消了知道者的一类动作本身，无论其后续作为多么自律。前者是可治理的偏，后者是不可逆的失；混同两者，就会把"换人"误读为一种过度谨慎。

五、概念与命题：盲遍、预走与判据

现在把判据立起来。

定义一（盲遍）：设某人面对一个对其而言的具体未知。在该未知的答案对其不可得的状态下，他把自己的判断交出并落到可核形式（写下、提交、当众宣布、签名）——这一事件称为一次盲遍。

定义二（预走）：若答案在承诺落定之前对当事人可得，则该件交付上的盲遍被取消，称该件为预走。

定义三（门位次）：一件交付的门位次 k ，是答案到达时刻相对承诺落定时刻的次序读数：答案后到， k 取正；答案先到， k 记零。其清点规程见第七节。

在三个定义之上，本文的主张压缩为三条命题。**命题一（唯一性）**：对每一具体未知，每人恰有一遍盲遍——承诺落定或答案先到，二者任一发生，该遍即告消耗。**命题二（不可赎回）**：被预走的盲遍不存在任何本人可执行的恢复程序；系统层面的补救只有换人与换题，而两者都不归还原件。**命题三（位次化）**：当答案供给的边际成本趋零、默认到达位次整体前移时，毕业生之间可持久的竞争差向门位次的分布集中——档案内容的区分力单调让位于形成时序的区分力。命题一、二由第四节的反例与第六节的结构特征承担论证；命题三是本文对题面的正面回答，其可证伪形式见第十六节。

三重否定式：人工智能时代大学毕业生的核心竞争力，不是任何可显露的成品品质——因为品质的读数本身随判读者与制作者的知情状态改变；不是任何可补修的养成路径——因为路径里有一段只在无知态存在，错过无处补；也不是任何随人工智能一同丰裕的条件——因为答案的丰裕正是耗散它的力。核心竞争力是他名下仍在发生的盲遍：答案到场之前完成的可核承诺，以及把答案排到承诺之后的那份位次安排。

判据（零情态词）：取任一件交付，核对两个时刻——该问题的答案对当事人可得的时刻，与当事人承诺落定的时刻。答案后到，盲遍成立；答案先到，该件交付上不存在判断的形成，无论其档案与自走的档案如何逐字相同。

这里必须立刻说清三个"不是"。第一，盲遍不是无知崇拜。命题不说不知道好，说的是：判断只在"不知道且交付"的合取里成形，缺一不可；光不知道不交付，什么也不发生。第二，盲遍不是反对使用人工智能。答案在承诺之后到场，是最好的老师——判据只认位次，不认用量。第三，盲遍不是主张挣扎更有效。有效与否是程度问题，可以争论、可以测量、可以折中；本文的主张是存在问题——答案先到，那一遍不是变差了，是没有了。这两种主张会在第十二节正面交锋。第四，盲遍不是怀旧论。它不主张回到没有搜索引擎的年代——检索一条既有事实不消耗任何盲遍，因为"世界上已知而我未知"的事实性知识本来就不该由我重新发明；被预走的只能是那类需要我作承诺的判断件。把查资料与要答案混为一谈，是对本文最省事的误读。

六、预走的三个结构特征

预走的第一个特征是**无声**。谜底先到的那一刻，当事人的主观体验是效率：更快懂了，少走弯路了。没有任何一个环节会报警，因为从信息的账上看，确实什么都没少——少的那样东西从来不在信息的账上。

第二个特征是**不损档案**。预走产出的记录与自走产出的记录逐字不可分。这不是取证技术暂时落后，是结构性的：档案登记的是承诺的内容，而预走取消的是承诺的性质；内容可以相同到每一个标点。第一节那两份档案的不可区分，根源在此。

第三个特征是**不可赎回**。判定一道门是不是单向门，有一个便宜的试法：出钱能不能买回来？能买回来的是成本，不是单向门。补课买得回落下的进度，重修买得回学分，唯独买不回“对这道题的不知道”。可以换人——找一个未被剧透的人替你猜——但换人是系统层面的绕行，不是个人层面的赎回；而当答案的供给端变成全天候、嵌入式、默认开启的人工智能时，连系统层面的绕行也开始失灵：未被剧透的人本身成了稀缺品。盲法研究对此早有制度自觉——被污染的判读者只能换掉，不能洗净——只是它把这件事记在“偏倚防控”的账上，没记在“一种资源被永久花掉”的账上。同一份材料，两种记法，评价相反：它说这是要防的污染，本文说这是被花掉的、恰有一遍的资源。

换人这条绕行路值得再走深一步，因为它藏着本题的时代之变。单向门并不是新事物——谜语自古只能猜一次——但它此前从不构成社会问题，原因是换人便宜：一个判读者被污染，试验换一个；一道题被某届学生传开答案，考试换一道；总有足够多“还不知道的人”与“还没被做过的题”。整个评价与培养体系其实一直靠这个隐形的供给池运转，只是从没给它记过账。人工智能改变的正是这口池子的水位：答案供给第一次做到全天候、嵌入式、对全人口同时开启——它不是污染某一个判读者，是同时预走一整代人对同一批题的第一遍。换人失灵不是因为门变严了，是因为门外排队的“未知情者”没有了。这就是为什么“核心竞争力将是什么”这道题必须在此刻重问：以前盲遍从不稀缺，谈不上竞争；此后它是稀缺品里最特殊的一种——不可储存、不可转让、按人按题各恰一份。

七、测量：门位次、盲遍率与清点规程

盲遍的旗舰读数是**门位次**，记作 k ：在一件交付的形成序列里，答案到达落在承诺落定之前第几步还是之后第几步。约定：答案在承诺之后到达， k 取正，盲遍成立；答案在承诺之前到达， k 记零，该件为预走。

两条随行读数。**盲遍率 b** ：一个人最近 N 件同域交付中 k 为正的份额。**二次落差 Δ** ：在迁移新题上，预走者与自走者后续表现之比——它承接一条老实验传统的读数形状：第二次做同类事的表现除以第一次，落差在哪一步出现，单向门就在哪一步。

清点手册七条，全篇一个口径。其一，承诺以可核形式落定为准：写下、提交、口头宣布于他人、签名，皆算；心里想过不算。其二，答案可得以“当事人一步之内可取得”为准：解析印在题旁、补全悬在光标上、答案在同屏聊天窗里，皆为可得；世界上存在答案但当事人取用需一次真实检索的，不算先到。其三，部分答案按最小承重件计：提示到了思路没到， k 按思路一步记；思路到了， k 记零。其四，同域指判断类型同族，不指题面相似。其五，分母只收有承诺记录的交付；没有承诺这件事的时段，不进任何分母——本文丈量的是承诺与答案的先后，不是承诺的有无。其六， b 只作场内比较，不跨域折算。其七，读数指位置，不指人： k 与 b 描述一件交付、一段履历里答案到达的位次结构，不给任何人贴“预走者”的身份标签。

给一个最小算例。某学生周一到周五交了五件作业：周一那道证明题，他先写满两页草稿、按下提交，再看解析——答案后到， k 为正。周二的编程题，编辑器补全在他构思之前

弹出了整段实现，他读懂后采纳——答案先到，k 记零。周三他卡住，向模型要了"思路提示"但没要代码，随后自己写完提交——按最小承重件计，思路这一步被预走，实现这一步是他的：k 按实现一步记正，但记录里注明思路件先到。周四事实性检索（某个接口的参数顺序），不进分母。周五他先提交再对答案，k 为正。本周分母四件，k 为正者三件，盲速率 b 为四分之三。整套清点五分钟可完成，材料全部来自他自己的提交时间戳与工具日志。

八、形成序列模型：轮次时间轴与关键反事实

判据要立得住，必须把那道门钉在序列的具体位置上，并且回答一个反事实。图1把一件交付的一轮走完：

图1 交付形成的轮次时间轴（含第⑦步反事实）

第 n 轮

- ① 未知到场（一道题、一个病人、一段要写的程序）
- ② 当事人调集已有之物，形成一个待交的判断
- ③ 承诺落定 ←—— 单向门在这一步：从此对这一未知，
| "无答案在场的承诺"这一动作不再存在
- ④ 答案（判卷、结局、运行结果、模型的版本）到场
- ⑤ 当事人以承诺对照答案，差处成为下一轮的材料

第 n+1 轮

- ⑥ 面对新未知时，先动的是他自己的判断，还是先等答案
- ⑦ 反事实：若第 n 轮里④先于③（答案先到），⑥会不会不同

门位次 k 丈量的就是③与④的先后。第⑦步是全篇的生死线：如果④③次序颠倒而⑥毫无不同——预走的人在下一轮照样先动自己的判断、后续盲承诺表现与自走者无差——那么"先后有关"就只是常识水位的相关，本文整个位次层作废。第十六节的证伪条款F一与F四，就是把第⑦步做成可跑的实验：随机安排答案到达的位次，看迁移新题上的后续差异。现有旁证指向次序颠倒确有后果：向高中生开放不设防的答案供给，撤走之后，他们在无辅助考试上比从未用过的同伴差约百分之十七（Bastani et al., 2024）；而同一模型改为扣住答案只给引导，损害消失。但旁证是旁证，它测的是成绩不是动作性质，本文不拿它当已完成的检验。

归格声明：本文回答的是一道"是什么"的题，而它交出的答案落在路径一侧——核心竞争力在这里不是一件显露出来的东西，也不是一片条件，而是一段路的存续：那段只有在答案缺席时才走得成的路，正在被"答案已在场"与"答案供给无限"两件事合力取消。就本文所知，对这道题的既有讨论均把答案落在显露一侧（某种品质、某种残差、某种判别力），落在路径一侧的，本文是第一份，因此判决性的分界押在第⑦步上：凡显露侧的说法，对"答案到达位次不同而档案相同"的两人不作区分预测；本文预测两人后续盲承诺表现有差。这不是指责显露侧读法粗疏，而是指出它的量程边界：一切从成品读人的方法——考试、作品、测评、履历——共享一个几何限制，它们读的是③号点位之后留下的东西，而门位次是③与④之间的次序事实，成品形成之时次序已经坍缩进内容里，如同冲洗出来的照片上读不出快门与闪光谁先谁后。要读次序，只能在次序还活着的时候读：形成期的时间戳，或者现场铸造一个新的③。

九、履历类型学与靶格：全对而未走过

用两条轴把毕业生的履历切成四格：横轴是答案到达位次（承诺前／承诺后），纵轴是档案完备度（齐／缺），得表1。

表1 履历的四格类型学

	档案齐全	档案残缺
答案后到（k 为正）	①自走成档：判断在成形，且留了痕	②自走失记：走过了，没记下——吃亏在读法，不在实质
答案先到（k 记零）	③全对而未走过：每题都对，每题都是答案先到	④双空格：既未走过也无档案

靶格是③。它有三个签名，合起来解释了为什么这一格在现行读法下完全隐形。签名一，当期指标改善：答案先到确实提高当期正确率，所以一切以当期正确率为准的评价都会把③格读成优等。签名二，失效无信号：③格的档案与①格逐字不可分，任何审计翻不出差别。签名三，越正规化越严重：教辅把解析印在题旁，题库把答案键挂在提交页，编辑器把补全嵌进光标，每一次"更好的学习支持"都把答案的默认到达位次向前挪一步——把学习支持做得越制度化、越体贴，③格的产量越大。一项被广泛讨论的脑电研究（Kosmyna et al., 2025）观察到，长期由模型代笔的写作者在换回自笔时神经参与度持续偏低，作者称之为认知负债；无论其测量还能再争多久，它至少说明③格开始在档案之外的地方留下可测的影子。

第一节那两位毕业生，现在可以定位了：第一位在①格，第二位在③格。四年逐字相同的档案，是①与③的不可分——而人工智能时代的结构性事实是：答案供给的默认位次整体前移，③格正在成为默认产线。

③格不是思想实验，它今天就有产线，举三处。其一，编程教育。代码补全工具的默认形态是"建议先于构思"：光标停顿半秒，整段实现浮现。学生采纳、运行、通过，提交记录完美——而那道题上"无实现现场时构想实现"的一遍已被取消。教师群体近来反复报告同一件事：学生交得出跑得通的代码，被问到"为什么这个方案而不是那个"时当场失语；有课程据此把评分从"交付了什么"改成"能否解释为什么"，这等于承认档案已读不出①③之别，只好现场铸造小型未知。其二，数学教辅。解析与题目同页印刷是市场的默认版式，因为家长按"孩子当天能否做对"付费——签名一（当期指标改善）直接变现，于是③格由市场机制批量生产。其三，语言与写作。模型代笔的初稿先到，学生在其上修改——修改是真实劳动，档案更厚了，而"面对空白页作出第一个结构承诺"的那一遍没有发生。三处产线共享同一个结构：每一处的答案前移都以"更好的支持"为名，每一处的当期指标都真实改善，每一处的取消都不留痕。

十、可走性的再生：春化研究的机制证词

命题走到这里，最自然的反驳是绝望论：若盲遍不可赎回，而人工智能把答案的默认位次一路前移，那么这一代学习者岂不是被判定为集体预走、无路可救？植物学里有一件事直接顶住这个推论。

许多温带植物必须亲历一冬低温才能开花，这个过程叫春化。它是严格的亲历件：低温必须由这株植物自己、以足够的时长坐过，短了不算，别株代替不算——这门学科为此打过二十世纪最惨烈的一场学术战争，李森科一派主张春化态可以速成、可以遗传，被几十年的实验逐条埋葬。现代判据已经落到分子上：低温逐步沉默开花抑制基因FLC，沉默态经有丝分裂稳定维持，一株春化过的植物用组织培养再生出的新株可以不经冬天就开花（Burn et al., 1993）——记忆是真的，走过就是走过了。组织培养那条证据值得单独看一眼：把春化

过的植株取单个细胞培养再生，新长成的整株不需要再过冬——“走过冬天”这件事写进了细胞的可继承状态里，连推倒重建一具身体都抹不掉它。亲历件能硬到这个程度，反过来解释了李森科一派当年为什么必须造假：他们承诺的是“走过可以速成、可以代继”，而这两条恰恰是这个系统用分子机器专门封死的。

但这门学科最锋利的一条不在记忆，在**重置**：春化态在有性生殖时被主动抹除（Sheldon et al., 2008），FLC在种子发育完成前被重新激活，保证下一代必须重新过冬。这不是记忆机器的失灵，是一台专门的酶学装置在执行抹除——找到那个去甲基化酶ELF6的工作（Crevillén et al., 2014）发表在二〇一四年：把它突变掉，后代就“继承”了半个春化态，不用亲历冬天也能开花。换句话说，**演化专门花成本造了一台重置机，防止“走过”被继承，以保证每一代的可走性。**

这条对本文有两个后果。第一，它把盲遍从悲观命题翻成设计命题：不可赎回说的是同一人对同一未知；而未知的供给是可以被制度性再造的——出题者的职业本质就是造未知，未解问题池、不带解析的题库、答案延迟发布的作业系统，都是人类版的ELF6。可走性不可赎回，但可再生。再生是有账单的：一道不配解析的好题、一个答案延迟发布的作业系统、一位肯憋住不说的老师，成本都实打实高于把答案印在旁边——市场自发的方向永远是答案前移，因为前移的收益（当期指标）归付费者即期兑现，前移的代价（形成停摆）由学习者远期承担，中间隔着四年到十年的账期。所以未知的再造不会由市场自动供给，它像基础研究一样，是一件必须被制度性买单的公共品——这也是“大学最后的不可替代之处”这句话的经济学脚注。第二，它给“核心竞争力”补上制度的一半：一个毕业生的盲遍率不只由他自己的纪律决定，也由他所处环境的答案默认位次决定——大学与雇主握着重置机的开关。

十一、初遇的锚定：免疫印记的推论

还有一门学科从相反的方向逼近同一结构。免疫学自弗朗西斯（Francis, 1960）提出“原始抗原之罪”起就知道：免疫系统没有第二次初遇。你一生中遇到的第一株流感病毒，终身规定你此后对整个流感家族的应答形状——后来的每一次感染与接种，都被解读为对第一株的变奏。发表于《科学》的分析（Gostic et al., 2016）给出可称壮观的证据：以一九六八年病毒株更替为界，不同出生年份的人群对H5N1与H7N9两种禽流感表现出方向相反的重症保护格局——你几岁那年流行什么，决定了你六十岁时怕什么。

免疫学对第一遍的处置与盲法相反：盲法**护**着第一遍（把知情挡在判读之后），免疫系统**花**掉第一遍——第一次初遇不可选择、不可推迟、花掉就花掉了，而且花在哪里就终身锚在哪里。疫苗接种的全部智慧，可以读成对这个不可选择性的位次管理：既然第一遍必然被花，那就抢在野毒之前，用一个设计过的、代价低的版本去花它。

这给本文加了一条毕业生层面的推论：盲遍不仅有“有没有走”的问题，还有**花在哪里**的问题。四年本科大约能承载的高质量盲遍是有限的——每一遍都要真实的未知、真实的承诺与真实的对照，这三样都贵。把它们花在什么域上，近似不可逆地规定此后判断力的形状。这不是“专业选择很重要”的老话：老话说的是知识存量的方向，这里说的是**第一遍的主权**——同一批知识可以后补，第一遍花掉的位置不能重选。

由此得两条教育推论。其一，低龄段的答案供给比高龄段危险一个量级：印记的锚定力随初遇的早晚递减，给八岁孩子的作业本旁边印解析，与给二十八岁的工程师配补全工具，在本文的账上不是同一件事的两个剂量，是两件事。其二，“用便宜的版本抢先花掉第一遍”这条疫苗逻辑，可以整段搬进教学设计：既然学生对某类问题的第一遍必然被花，那就由教师设计一个低代价、高对照质量的场合去花它——课堂上的首次独立尝试，本质上就是认知的减毒活疫苗。

十二、理论定位：与最近的几种说法的分界

柏拉图的《美诺》离本文最近也最远。美诺悖论说：求知不可能——已知者无需求，未知者不识所求。柏拉图的解是回忆说：一切学习都是重走灵魂早已走过的路。本文与它正面对：恰恰存在一段不可重走的路，而学习的实质就发生在那一段上。判定形式：回忆说预测“被提前告知”与“自行想起”殊途同归（《美诺》里那几个几何奴隶无论如何都会抵达）；本文预测两者在后续迁移上分开。细看那段对话会发现苏格拉底自己守着本文的纪律：他向奴隶只发问、不给答案，让每一个错误答案先落定、再被几何事实驳回——奴隶两次答错、两次亲眼看着自己的答案被对角线否定，然后才抵达。柏拉图用这场示范论证回忆说，而这场示范的每一步操作，恰恰是在保护盲遍：假如苏格拉底开口直接给出对角线，按回忆说毫无损失（唤醒即可），按本文则整场学习不曾发生。理论与示范打架，本文站在示范一边。

赫拉克利特说人不能两次踏入同一条河。分界一句话：他说河变了，本文说踏河者的知识态变了——把河冻住（题目一字不改），第二次踏入照样不是第一次。

后见之明偏差族（Fischhoff, 1975 ; Fischhoff & Beyth, 1975）是本文机制层最大的一位邻居，本节必须把账算清。菲施霍夫五十年前的实验给出的正是一个不可逆的知识态：得知结果之后，人无法恢复自己此前的不确定状态，他称之为潜行的决定论——过去在记忆里被围绕已知结果重建成一条必然的因果链，而重建是无意识的。五十年综述（Fischhoff, 2025）报告了两条对本文极其有利的实况：单纯警告受试者存在此偏差没有可察觉的效果；各类去偏尝试大多失败，最有希望的一条是帮助人们重建当时的视角——即换回一个尚未知情的立场，而那正是本文所说的换人或换题在个体内的微缩版本。

分界有三处，逐条写明。**其一是账目的位置**：后见之明记的是回溯账——已知结果如何扭曲了对过去的判断（我早就知道会这样）；本文记的是前瞻账——答案先到如何取消了尚未发生的那次承诺。前者的对象是记忆，后者的对象是一次本可发生而未发生的动作。**其二是资源观**：后见之明把知情态当成一种持续作用的扭曲力，因而处方是持续对抗（记录、复盘、重建视角）；本文把盲遍当成一份恰有一遍、一次性消耗的份额，因而处方不是对抗而是排序。**其三，也是最要紧的一处，本文必须诚实申报**：后见之明文文献的标准处方——在结果揭晓前把预测写下来，即决策卫生与事前验尸——在操作层面几乎逐字覆盖了本文第十七节的旗舰形态一。相同的动作，两套账：在那边，写下预测是**去偏工具**，其价值在于事后可以拿它对照、防止记忆篡改；在这边，写下预测**就是盲遍本身**，其价值不在事后对照，而在这一次判断确实被生产出来了。可分实验因此存在：取一批预测日志，事后一律不作对照、不作复盘——按去偏论，未经对照的日志几乎无用；按本文，它已经完成了它全部的工作。

知识的诅咒（Camerer, Loewenstein & Weber, 1989）从市场实验一侧给出同族证据：拥有信息者无法准确模拟不拥有该信息者的状态，且这种无法不因激励而消失。它与本文共享不可逆的知识态，分界在承受者：诅咒说的是知情者对**他人**的模拟失准（因而是沟通与定价问题），本文说的是知情者对**自己**的一类动作的丧失（因而是形成问题）。

变革性经验（Paul, 2014）是这条线在哲学上的正主：某些经验在你亲历之前无法被认识论上模拟，一经亲历则你与你的偏好都已改变，故此类决策不能靠事先想象来作。它与本文同构于单向门，分界在门的宽度与频次：变革性经验讨论的是稀有的、改变整个人的大门（生育、移民、失明），每人一生数扇；本文讨论的是稠密的、按题计数的小门，每人每天数扇——而正因为小、因为多、因为无声，它才可以被产业化地批量关闭，这恰是本文的题域。

合意困难与前测效应（Bjork, 1994 ; Kornell, Hays & Bjork, 2009）是教育心理学一侧最该被本文认账的仪器：先猜后学优于直接学，哪怕猜必错——这条实验范式里，答案的到达次序早就是一个被随机操纵的自变量。本文对此不作任何仪器上的原创声明：F—实验

若被执行，用的就是这套现成范式。差别只在读什么——前测效应读的是同一批材料的记忆保持率，本文读的是迁移件上无脚本承诺的表现；仪器共用，因变量不同。

《论语·述而》的“不愤不启，不悱不发”是本命题的汉语经典层版本，且措辞比现代文献更贴：启与发是教师的动作，愤与悱是学生自己先撞上未知而未通的状态；孔子给的是次序纪律——学生的那一遍必须先发生，教师才允许开口。本文与它的分界在于，它交出的是师道规范（该不该），本文交出的是存在判据与读数（是不是、第几步、多少份额）；两千五百年的规范之所以直到今天才需要读数，是因为在此之前从没有一台机器能把开口这件事做到全天候、零成本、默认开启。

生成效应（Slamecka & Graf, 1978）实证了自己生成的信息记得更牢。它是本文最老的实证近亲，但它记的是记忆账——生成比给予多百分之几的保持率；本文记的是存在账——不生成而先给予，被取消的不是保持率，是那个动作。剂量与存在，两本账可以同向，也可以在保持率无差时依然分开。

香农的一次性密钥（Shannon, 1949）在形式层上与盲遍同构：一次性密钥的安全性来自单次使用，钥匙重用一次，整个体系即告破——冷战期间的维诺那破译就是靠对方重用了本应只用一次的密钥。把内容全部换掉、只留结构，两边剩下同一句话：这类对象的全部价值系于“恰有一遍”。分界在于密钥的单次性是设计出来的约定，盲遍的单次性是认知动作的构成性事实——前者可以印一本新密钥簿，后者只能换一个人。

测验学的题目曝光控制（Sympson & Hetter, 1985）是方法学内部又一位近邻：计算机化自适应考试为每道题设曝光上限，因为一道题被考生群体见过之后，它测量的东西就变了——题库管理的日常，就是守护每道题对每位考生的“第一次”。同一结构，账记在相反的一侧：曝光控制守的是题的判别力（题坏了，换题），本文守的是人的形成（人这一遍没了，换不了人）；两者在考场相遇——泄题之所以是双重事故，因为它同时毁掉一道题和一届人的同一遍。

有效失败（Sinha & Kapur, 2021）与**辅助两难**（Koedinger & Aleven, 2007）是本文在教育研究内部最近的两位对手。前者以逾万人的元分析证明：先自行尝试、后接受讲解的学生，成绩胜过先听讲解的学生；后者把“何时给帮助、给多少”刻画为可优化的两难。两者与本文共享全部经验直觉，分界只有一条但足以判决：**它们把答案先到当效率问题——剂量不当、时机欠佳、可以调优；本文当存在问题——答案先到的那一件上，不是学得差了，是那一遍没有了。**同一份材料，两种评价：他们说“先挣扎，因为效果更好”；本文说“先承诺，无所谓效果好坏，那是唯一的一遍”。可分实验：找一批“先讲解”与“先尝试”在保持率上无差的题型（元分析承认这类题型存在），本文预测两组在更远的迁移承诺行为上依然分开；效率论对此无预测。

“受保护的首次尝试”是二〇二六年出现的人工智能教学法框架，主张在六步教学环里保护两个时刻——首次独立尝试与最终无辅助检验——不许模型介入。这是离本文最近的当代邻居，近到本文必须感谢它：它已经把位次问题摆上了桌面。分界在理由：它保护初试，因为证据显示这样学得更好；本文保护初试，因为那一步之后没有第二次。理由之差决定政策之差：按效果论，当某天出现“答案先到也能学得一样好”的新证据，保护即可撤销；按存在论，那种证据不可能出现在“同一件”上，只能出现在迁移件上——保护不随效果证据涨落。

认知负债与代笔研究（含前引脑电研究，及“生成式人工智能会损害学习”的高中实地实验）记录了依赖答案供给的亏损。分界：亏损账默认可偿——补练、戒断、再教育；本文的存在账里有一部分不可偿。同一族里的《当答案成为廉价品》一文以“认知发生权”命名学习者对自身认知过程的主权，与本文最为亲近；分界在处方的形状：主权论的处方是守住一条回路——保留学习者参与的环节即可；本文的判据更窄也更硬：回路里若答案仍先于承

诺到达，参与照样是预走，守住回路不等于守住位次。判定形式写死：若未来出现任何不经换人的“复盲”程序——能把已知情者恢复为可对同一未知作盲承诺的状态——本文即伪。

十三、与同批诸文的分界

同一道题，本文所属的同批工作已给出六个答案，锚位各不相同，必须逐一划清。《现配优势》把竞争差锚在人与场之间：优势由平庸构件加一手配组当场做出，读数是交接后的复现份额。《零判份》锚在产出与世界现存样本之间：一次合格交付里判别力为零的那一份。《脱稿步》锚在产物与作者的关系上，其最新版本把竞争力立为可续签的署名——署名是一份没有到期日却已经会到期的凭据，效力只能在自己产物的延长线上被当场追问时无脚本重做，读数为续签成立率。《常席差》锚在跨配组的残差上：人是唯一不变项时才读得出的那份符号一致。《回改》锚在时间回溯轴上：已完成的裁定被后到证据改写符号。《敞问》锚在自我轴上：只能由本人合上、一经外部裁定代合即失效的问。

本文与它们的总分界：六篇读的都是已经存在的东西在哪里显影，本文读的是一段路不存在。三条判决性对照。其一，对《脱稿步》：它的续签读一次被发起的追问，发生在交付之后，是对既有署名的当场重做；门位次读答案到达的时序，是形成期事实——续签检验一份判断力存不存在，盲遍解释这份判断力当初有没有被生产出来。两者一验一产，测量发起处隔着整个形成期。档案与追问成立率完全相同的两人，若形成期的门位次分布不同，本文预测其后续盲承诺表现不同——《脱稿步》对此不作预测。其二，对《敞问》：它预测外部裁定使那个问失效，登记即毁；本文预测零裁定、纯知情即毁——可分实验是不登记、不裁定、只把答案悄悄递到承诺之前。其三，对《现配优势》与《常席差》：两者在“人不变、场轮换”的横截面上读数；本文在“答案何时到”的纵剖面上读数。同一个人可以现配优势极高而盲遍率为零——他当下的配组才华全部来自毕业前攒下的判断力存量，而存量正在停止增长；两套读数给出相反的五年期预测：横截面读数预测他持续领先，本文预测他的领先以形成停摆为代价，在域迁移到来时一次性兑付。对《零判份》与《回改》再各补一句。《零判份》问一次交付里有多大部分额判别不出作者水平，它读的是产出与世界的关系，本文读的是产出与其形成过程的关系——两者可以同向叠加：预走产线批量制造的，正是零判份极高的合格品。《回改》说裁定不终局、后到证据能改写符号，本文与它在时间轴上恰成一对镜像：它管承诺之后答案还会怎样回来改账，本文管答案抢在承诺之前会取消什么——一个读后事，一个读先后。

十四、两次事先登记的检验

本文按事先登记的方式检验了自己最便宜的一条可证伪主张：第四节那个“另一种动作”层。先说为什么押这一层。全篇的推理链是：另一种动作，故不可校正；不可校正，故只能换人；换人失灵，故位次成为竞争差。链条最上游、也最便宜可查的就是第一环——它若倒，后面全塌，而查它只需要已发表的元研究统计量。若知情只是一个可校正的稳定偏移，各试验的夸大应方向高度一致、幅度同质，一个校正系数即可洗净——那样的话，盲遍就多余了，偏倚防控的老账足够。登记在读取任何统计量之前写死三件。判负条款甲：若逐试验的比值比之方向一致率不低于八成、且异质性统计量I方低于百分之四十，判“稳定偏移”成立，本主张负。判负条款乙：可提取的独立比较不足八项，判数据不足。另有不利结果处置与停止规则各一条，并申报了一处已入眼数据：合并比值比之零点六四为本文作者事先知晓，故判负条款只押作者事先不知的方向分布与异质性。

结果如下，照原样登记。异质性半边：姊妹研究（量表结局，十六试验）I方为百分之四十六且元回归无法解释；二〇二五年扩大更新在行业资助亚组内I方为百分之六十五，药物与非药物亚组方向相反，主观度、易感性、卷入度三条元回归全部空手。方向半边：二〇一二年的逐试验比值比之极差为零点〇二至十四点四，跨过一，但逐试验的方向计数未能

提取——原文全文两处访问被拒，本文如实登记这半边没有直接读数。按登记条款：判负未触发，"另一种动作"层获得支持，但因方向半边缺直接读数，本文把支持等级自降半档，称之为"对已发表统计量的事先登记再判读"，而非对原始数据的重跑。二〇二五年更新里三条元回归空手与亚组方向反转，列为事后线索，明标不在登记之内。这条检验欢迎任何人拿逐试验数据做完：数据在已发表的森林图里，条款在本节里，一个下午足够。

第二次登记：位次是否被记录

第一次检验查的是机制层（知情之后的判读是不是另一种动作）。第二次查的是命题三的一条制度可观测面：**若门位次真是竞争差之所在，那么现行可携带记录里应当找不到它。**这条推论便宜、可机械核验、而且——如果被推翻，本文第八节关于"形成期时序读不出"的判断就要收窄。

登记在读取任何标准文本之前写死，另加一条独立性申报：同批《脱稿步》曾审过同一机构的另一份标准（学习者综合记录 2.0），故本次主检对象改为**此前未审过的开放徽章 3.0 规范**（1EdTech，最终版，二〇二六年六月），以免两次检验共用同一份样本。三条判负条款：**甲**，若存在至少一个带类型（非自由文本）的字段，其语义为"学习者承诺/提交时刻与答案可得时刻的先后或间隔"，则推论判负；**乙**，若可审时间字段中有三成以上其触发方在持有者一侧（由学习者行为定义而非签发方设定），则"偏侧"主张判负；**丙**，若可审字段总数不足十个，判数据不足、本次检验不算完成。

审计对象为该规范的正式 JSON 模式文件，逐类逐字段清点。结果如下，照原样登记。

表2 开放徽章 3.0 的时间字段与触发方（全量七个）

字段	语义	触发方
awardedDate	授予时刻	签发方
validFrom	凭据生效	签发方
validUntil	有效期终点	签发方
activityStartDate	学习活动起点	由签发方记录
activityEndDate	学习活动终点	由签发方记录
proof.created	签名时刻	签发方
dateOfBirth	出生日期	第三方事实

计数：被审类二十二个，字段总数一百六十三个；其中时间/时序字段七个；触发方在持有者一侧者零个（零成）。**三条判负条款均未触发**：条款甲——七个时间字段无一表达承诺与答案的先后，最接近的两个（活动起止）记的是**活动时段**，与本文所问的**位次**是两回事，一个人在同一段活动时间里，答案可以先到也可以后到，这两个字段对此完全沉默；条款乙——零成，远低于三成阈值；条款丙——一百六十三个字段，样本充足。推论 P 获支持。

有一处必须诚实交代，它削弱本文的强表述。该规范设有背书自由文本插槽，理论上任何人可以在那里写一句"此人在看到参考答案之前已提交方案"。所以"完全没有地方记"是错的，当场划掉：**能记，只是记成一段没有类型、不可比较、不可加权、不可查询的散文**。本文的主张据此收窄为：**没有带类型的位次字段**——而没有类型就没有排序，没有排序就进不了任何筛选流程，这正是第九节签名二（失效无信号）在制度层的对应物。

两次检验合起来给出的图景是：知情之后的判读是另一种动作（第一次，弱支持），而这个“另一种”在现行记录体系里没有位置（第二次，支持）。两份标准来自同一机构但属不同规范、不同数据模型，构成两个独立样本，结论一致。

十五、边界条件与反向证据

三件事真实地削弱了本文，照实写。

其一，二〇二〇年的大规模元流行病学研究（Moustgaard et al., 2020；跨一千余项试验）发现：平均而言，盲态与非盲态试验的效应差别不大。这一刀砍掉了“凡知情皆改判”的强版本。本文据此收窄：盲遍只在**判读带解释余地的域**承重——判断、设计、诊断、写作、策略，凡结论需要一个人的裁量去合拢的地方；客观结局的域（死了没死、编译过没过、称重读数）盲遍不承重，那里答案先到后到，承诺是同一种动作。

其二，心理学里的剧透研究（Leavitt & Christenfeld, 2011）发现：被剧透的读者对故事的享受平均不降反升。这一刀划开了两本账：盲遍管判断的形成，不管体验的品质。被剧透的电影照样好看，甚至更好看；只是“猜凶手”这个动作没有了。本文第二节的日常锚，据此收口在猜谜类认知动作上。

其三，迁移网络的稠密性。预走一道题不等于失去一个领域——同族的新题可以再造未知，而且人工智能自己就是历史上最高产的新题制造机：它花盲遍，也铸盲遍。这一刀最重，它逼着本文把重心从存量挪到位次：真正的竞争差不在“还剩多少没被剧透的题”（那个池子可再生），而在**答案到达的默认位次落在哪里**——一个把答案默认排在承诺之后的人与制度，在同样的答案丰裕里持续成形；反之则持续预走。第十节的重置机之所以是出口的一半，正因如此。

不适用边界还有两条。其一，纯操演技能（打字速度、乐器指法）：那里的形成靠重复而非承诺，盲遍读数无定义。其二，答案本就不存在的开放域（前沿研究、创业）：那里人人盲遍率为一，读数失去区分度——但这恰好反过来解释了一件事：为什么最像前沿的训练最出判断力。其三，协作域的归属：多人共同承诺的交付，盲遍按承诺主体记在组上，组内个人的份额本文不载——那是同批《现配优势》与《常席差》的地盘，两套读数在此交接而不重叠。

十六、可证伪性与判决性预测

四条证伪条款，全部落在可执行的第二档。F一：随机化答案到达位次的同质多案例实验——同一批学习者、同族题目，随机安排答案先于或后于承诺到达，各不少于六个案例、事先登记；若两组在迁移新题上的后续表现无差，本文位次层伪。F一是本文唯一能把位次层从假说抬成读数的实验，本节把它写成可直接执行的协议，任何一门课都能跑，不需要本文作者参与。**设计**：同一批学习者、同族题目，随机分派到两种到达次序——甲组每题先提交自己的判断再获得参考解，乙组每题先获得参考解再提交，题目、答案、讲解、总用时四项完全相同，**唯一自变量是次序**。**样本**：不少于六个独立案例（班级或题组），事先登记假设、分析方案与停止规则。**关键控制**：总信息量必须严格相等——这样任何组间差异都只能记在位次头上，记不到内容或剂量头上；这一条不做到，实验作废。**结局指标**：必须取**迁移件**上的无脚本承诺表现，不取原件——原件上的复述两组都会接近满分，那正是第九节的签名一，用它当指标等于自己制造零结果。**预登记的判负条款**：若两组在迁移件上的表现差异的置信区间包含零，且该结果在至少四个独立案例中重复，本文位次层判伪，无需辩解。**次级指标**：两组在迁移件上“承诺前主动检索答案”的比率——本文预测乙组显著更高，即预走会自我强化，形成本文所说的产线。**伦理**：两组最终获得完全相同的教学内容，差别只在同一堂课内的先后，不构成教学剥夺；第十九节三条伦理件全程适用。F二：若出现任何不经

换人的复盲程序，能把已知情者恢复为可对同一未知作盲承诺的状态，本文单向门层伪。F三：若两态判读之差可被一个跨域稳定的校正系数消除（低异质、单方向），本文“另一种动作”层伪——即第十四节的口径。F四：若预走件与自走件在“迁移件上的后续盲承诺表现”上无差、且该无差跨域稳定，则本文整个位次层伪，第八节第⑦步的反事实按“次序无关”结案。

赌注，写死日期：至二〇三一年十二月三十一日，主流学位授予与校园招聘的评审系统不会把“承诺时间戳早于答案可得时间戳”设为必填的结构化字段。若届时该字段已成必填标配，本文对制度惰性的判断错误，此节作废——那将是本文乐于输掉的一注。

十七、实践含义：毕业生、大学与雇主

对毕业生，本文只有一条可执行的话：把答案排到你的承诺之后。用人工智能，尽管用，但在按下“问”之前先写下你自己的版本——两分钟的草稿就够，落定即可。你的竞争力不在你比模型知道得多（你不会），在你名下还在发生的那些盲遍：它们是你判断力唯一的产地，而且每一遍恰有一次。这一条落成四个具体形态。形态一，两分钟前置承诺：任何要交给模型的问题，先用两分钟写下自己的答案版本——不求对，求落定；写完再问，模型的输出就从“替你走”变成“给你对照”。形态二，预测式阅读：读任何解法、论文、判例之前，先写一句“我预计它会怎么做”；这把被动阅读改造成一次微型盲遍，成本十秒。形态三，延迟核对：做完一批题不立刻对答案，攒到承诺全部落定再统一对照——把答案的到达位次从逐题后移到逐批之后。形态四，位次日志：每周用第七节的口径清点一次自己的 b，不给别人看，只给自己看——这个读数一旦向外申报就会开始撒谎（第十八节），向内使用时它是全篇唯一能自我校准的仪表。四个形态的共同点：没有一个要求少用人工智能，全部只动一件事——先后。

对大学，本文的推论是角色迁移：当答案供给不再稀缺，大学最后的不可替代之处是**未知与管位次**——出没有解析的题，延迟答案的发布，把“承诺先于答案”做进作业系统的时序里，像春化的重置酶那样，专职保证每一代的可走性。一所把学习支持做到答案永远先到的大学，是在用最体贴的方式取消自己的学生。落到操作上是三件事。第一件，时序化的作业系统：提交通道先于解析通道开放，承诺时间戳早于答案发布时间戳的，系统原样记录——不评分、不排名，只让时序成为档案的一部分。第二件，未解题池：每门课保留一批不配解析、不上题库、按学期轮换的题，专供盲遍——它们是这门课的种子库，教师的核心劳动从讲解答案转向培育未知。第三件，把“现场再走一步”制度化：答辩与考核的重心从“你交了什么”移到“在你交付物的延长线上，当场再走一步给我看”——这件事同批的《脱稿步》已给出完整读数，本文只补一句它缺的理由：现场追问之所以读得出真伪，是因为它当场铸造了一个答案尚未到场的时刻。

对雇主，本文给一条读档案的新法：档案内容读不出①格与③格之差，但形成期的时序痕迹读得出一部分——提交历史里承诺与修订的先后、能否在自己交付物的延长线上当场再走一步。已经这样做的面试环节（白板、系统设计、现场重构）之所以有效，本文给出它们此前缺少的理由：它们不是在考知识，是在现场铸造一个小型未知、强迫一次盲遍，然后直接读它。高校近两年回潮的互动式口试走的是近旁另一条路——就学生自己已提交的作品实时追问，使外包难以为继；两者相邻而不同伴：口试在既有产物的延长线上验证作者身份，本文的现场铸造在一个新未知上强迫形成——前者读的是《脱稿步》的续签，后者读的是一次新的盲遍，一套考核里两者可以并置，各测各的。由此推出雇主侧的两条操作与一条禁令。操作一：把面试题库当一次性密钥管理——每题用过即换，因为泄出的题对整个候选人群体已是“答案先到”，此时它筛出的不再是判断而是情报网络。操作二：入职后的头两年，为新人保留一条“先交方案再看基线”的工作次序——那是雇主唯一能亲手补铸盲遍的窗口，

过了判断成形的窗口期，给再多项目也只是在既成的判断力上加载荷。禁令：不得倒推——不得因为一位候选人的历史档案“答案可得性存疑”而降格录用；第七节口径第七条在此适用，读数指位置不指人，且历史时序多数不可考，可考的只有你现场铸造的那一次。

十八、讨论：三个理论后果

若命题一至三成立，至少三处既有理论需要改动，本节逐一写出改动的形状。

第一处在人力资本理论。存量框架把教育投资折算成可携带的资本，第三节已述其清单形态；针对存量框架的既有批评多半来自流量方向——能力会折旧，须持续再投资。本文加的不是又一次折旧提醒，而是一根缺失的轴：资本核算只登记“拥有什么、拥有多少”，从不登记“取得次序”；而按命题一，两份账面完全相同的认知资本，可以因形成期答案到达位次的不同而在迁移场合分道扬镳。这意味着人力资本的核算单元需要从（内容，数量）二元组扩成（内容，数量，位次）三元组——第三个分量不可由前两个折算，正如照片的像素矩阵折算不出快门与闪光的先后。

第二处在测量与评价理论。经典测量理论用信度与效度刻画一个测验的品质，两者都把测验当作对被测者无损的读取。第十二节的题目曝光文献已经打开了半扇门：题会被读坏。本文推开另外半扇：人也会被读坏——更准确地说，被答案先到的“讲评”读坏。由此需要一个新的品质维度，本文称之为形成核算：一个测评系统除了报告它测得准不准，还应报告它的答案供给史对被测群体的门位次分布做了什么。一份年年泄题却信度极高的考试，在经典框架里无可指摘，在形成核算下是一台预走机器。

第三处在人机分工理论。互补传统把分工当静态比较优势问题：模型做甲类任务，人做乙类任务。这个框架默认人对乙类任务的判断力是给定禀赋，而按本文，它是被持续生产出来的——生产线就是乙类任务上的盲遍序列，而模型的默认接入方式（补全先于构思、初稿先于空白页）直接设定这条生产线的开工率。于是分工设计在时间维上多出一个自由度：同一套人机配置，答案排在承诺之后是判断力的孵化器，排在承诺之前是判断力的休止符。第三节末尾那两项现场实验的普惠读法与本文读法之争，最终将由这条时间维裁决：若新手的即期获益伴随其独立判断在后续迁移件上的相对停滞，位次读法胜；若无此代价，本文第十六节的证伪条款自动执行。

十九、伦理与证据边界

本文的命题天然配着一类要在人身上执行的操作——扣留答案、延迟告知——所以伦理件不可省。三条写死。其一，安全关键场合不得执行：医疗、驾驶、施工等一切错误即时伤人的域里，答案（正确做法）必须先到，那里要的是合规不是成形，盲遍训练只能在仿真环境里做，且事后必须完整告知并复盘。其二，读数指位置不是人：门位次与盲遍率描述交付的时序结构，不得用于给个人贴身份标签，不得据以惩罚使用辅助工具者——答案在承诺之后到场即为正当辅助，残障与辅助技术使用者的合理便利不受此读数评价。其三，凡已按此类读数做出不利安排的（评分、录用、晋升），必须公开其编码规则并提供复核通道。

反噬预言，明写：盲遍率一旦进入考核，最省事的达标法是表演——表演性地不查答案、伪造承诺时间戳，而盲态本身不可审计，无人能证明你落笔时心里没有答案。这不是臆测，反噬的雏形早有档案：贝克等人（Baker et al., 2004）就系统记录了智能辅导系统里的“钻系统空子”行为——学生连按提示键直奔底层答案、再回填正确选项，系统日志一片优良。当年的钻空对象是提示阶梯，把盲遍率做成指标之后，钻空对象会升级成时间戳本身。所以 b 作为考核指标会在制度化的那一刻自毁。本文看到的出口只有半个：制度侧把重置做实（未知的持续再造、时序的默认设置），个人侧把位次当作对自己而非对他人的纪律。另外半个出口——如何在不逼出表演的前提下核验盲态——本文看不到，照实写在这里。能写

下的只有一个方向性观察：核验难题的根源在于盲态是内部状态，而本文全部口径恰恰绕开了内部状态、只认外部时序——所以出路若存在，大概率不在测谎式的状态核验里，而在把时序基础设施做成默认：当承诺通道天然先于答案通道，表演的成本会高过老实。这是猜想，不是结论，按猜想的等级放在这里。

二十、局限、自反与结论

先立此存照四条局限。其一，第十四节第一次登记检验只完成了一半半径：异质性半边有直接读数，方向计数半边因原文访问受阻而缺直接读数，完整版对任何持有逐试验数据者是一个下午的工作量；第二次检验虽三条判负均未触发，但它验的是记录体系的现状，不能反过来证明位次真的决定竞争差——那仍要靠F一。其二，位次层（命题三）现阶段是有旁证的假说：第十六节的F一实验尚未有人执行，本文对第三节两项现场实验的另类读法在该实验完成之前与普惠读法平权。其三，清点规程第三条的最小承重件判定，在深度人机连续协作的交付上判定成本会显著升高，该口径的判者间一致性未经预实验。其四，本文的机制证词取自医学判读、植物发育与免疫三个域，教育域内直接的两态对照实证（同一学习者、同一题、知情态与盲态的成对判读）尚不存在，构造它本身就是一项未来研究。由此排出三件未来工作的优先序：F一随机化位次实验居首，时序化作业系统的田野研究居次，低龄段答案供给的剂量一位次研究居三。

再作自反。用本文自己的判据给本文定位。本文的形成过程里，第四节那批统计量对作者是“答案先到”的——合并比值比之比零点六四在设计检验之前已入眼，作者据此把判负条款挪到未入眼的量上，这是对预走的承认与绕行，不是豁免。更要紧的定位：本文通篇是可陈述物——定义、读数、条款——而它自己主张判断力只在盲遍里成形；所以**读懂本文不构成任何人的一个盲遍**，把本文背熟的人在本文的账上一无所获。本文能做的只有一件事：改变答案到达的位次安排，让读者的下一个未知晚一点被剧透。这个定位的具体后果是：本文不可能靠被阅读而兑现，只能靠被执行而兑现——而执行发生在本文之外、档案之外，恰如它所描述的那样，无声，且不可代走。还有一处自反必须写明：本文自己的写作过程由人工智能完成，它的每一步都有海量既有文本作为“先到的答案”——按本文自己的判据，这份文稿的形成过程里没有一次人类意义上的盲遍。本文不因此豁免，恰恰以此作证：一份在一切读法下合格的文稿，完全可以在“判断形成”的账上一无所有；读者眼前这份档案本身，就是③格的一个诚实标本。

第十六节的赌注到期日是二〇三一年十二月三十一日。在那之前，每一位读者手里都还握有一些没被剧透的题。别让答案先到。

附录：一份示例清点——把第七节的口径跑完一整周

以下用一个虚构但逐件可核的案例，把清点手册七条全部过一遍实弹，供任何想复算 b 的人对照。主角是计算机系大三学生周某，清点窗口为某周的十一件候选记录，材料来源为课程系统提交时间戳、编辑器遥测日志与他本人的问答工具历史。

件一，算法课证明题。草稿提交于周一晚九时十四分，课程解析于周二上午十时发布。答案后到，k 为正。件二，同课第二题。他先看了上届学长流传的答案影印件（可得性：同屏一步之内），再写提交。答案先到，k 记零。件三，操作系统实验。编辑器日志显示：他敲下函数签名后停顿，补全建议弹出整段实现，他按下采纳键。k 记零——注意这里不需要追问他“是否本来也会这样写”，口径只认可核的先后，不认反事实的心理状态，这正是判据能机械执行的原因。件四，同实验的第二个模块。日志显示补全弹出后他按了拒绝，自行写了不同的实现再运行。答案曾先到但未被采纳——按最小承重件计：他的实现承诺落定于一个“另一版答案已在场”的状态里，思路件已被污染（他看见了一种做法），实现件是他的。k 按实现一步记正，附注思路件先到。这是全周最难判的一件，也是口径第三条存在的理

由。件五，向模型询问某标准库函数的参数顺序。事实性检索，不进分母。件六，英语作文。他让模型先出了全文初稿，自己改写了六成句子。修改量再大，结构承诺是模型的：k 记零。件七，数学建模小组作业。三人共同讨论后落定方案，答案（往届优秀论文）在讨论之后才被调出对照。协作件，盲遍记在组上，k 为正，个人份额不裁。件八，他给学弟讲了一道自己上学期做过的题。不是他的未知，不进分母。件九，周五测验，闭卷。全场无答案可得，k 为正。件十，测验后对答案并订正。订正是对照动作，属于件九的第⑤步，不单独成件。件十一，他睡前刷到一条讲解视频，内容恰是下周作业的原题解法。这一件最值得停：他还没有对下周那道题作任何承诺，答案已经到了——下周那件交付尚未发生，已被预走。清点时它记为一次前置预走事件，待下周该件进入分母时直接判 k 记零。

结算：本周进入分母者七件（一、二、三、四、六、七、九），k 为正者四件，b 为七分之四。两条读法随附。第一，b 约五成七这个数本身不构成评价——第七节口径第六条，b 只作场内纵向比较：若周某上学期同口径的 b 是八成，这个数就在报告一件事：他的答案默认位次正在前移，而他的成绩单不会显示任何异常。第二，全部十一件的判定没有一处需要询问周某的主观感受，全部落在时间戳与日志的先后上——这份可机械执行性，是门位次与一切“学习投入度”“努力程度”类指标的分水岭。

最后声明本附录的边界：它演示口径，不演示监控。第十八节的伦理件在此重复一遍——这套清点的正当用法是周某自己算给自己看；任何把它变成他人考核的用法，都会在启用的那一刻逼出表演性时间戳，读数随即失效。仪表装在自己手上是罗盘，装在别人头上是测谎仪，而这台测谎仪测不了谎。

注释与人机分工声明

本文由人类研究者提出问题与裁定方向，由人工智能系统在其指令下完成文献检索、论证组织、检验执行与初稿写作；事先登记文件先于一切数据读取写定。文中对已发表研究的每一处数字引用均可循参考文献复核；第十四节如实登记了一处未能取得的读数。字数与版本：正文含附录机检汉字数二一二五三，二〇二六年九月二日初版，同日扩订为学术论文版（第二版），再订为第三版（补机制层邻居与第二次登记检验）。

参考文献

（按作者字母序；中文文献与同批稿列于末。）

Acemoglu D, Restrepo P. Automation and new tasks: how technology displaces and reinstates labor. *Journal of Economic Perspectives* 2019;33(2): 3-30.

Arrow KJ. Higher education as a filter. *Journal of Public Economics* 1973;2:193-216.

Autor DH. Why are there still so many jobs? The history and future of workplace automation. *Journal of Economic Perspectives* 2015;29(3):3-30.

Baker RS, Corbett AT, Koedinger KR. Detecting student misuse of intelligent tutoring systems. *Proceedings of Intelligent Tutoring Systems* 2004:531-540.

Bastani H, Bastani O, Sungu A, et al. Generative AI can harm learning. Working paper, 2024.

Becker GS. *Human Capital: A Theoretical and Empirical Analysis*. New York: Columbia University Press, 1964.

Brynjolfsson E, Li D, Raymond LR. Generative AI at work. NBER Working Paper 31161, 2023.

Burn JE, Bagnall DJ, Metzger JD, et al. DNA methylation, vernalization, and the initiation of flowering. *PNAS* 1993;90:287-291.

Chouard P. Vernalization and its relations to dormancy. *Annual Review of Plant Physiology* 1960;11:191-238.

Crevillén P, Yang H, Cui X, et al. Epigenetic reprogramming that prevents transgenerational inheritance of the vernalized state. *Nature* 2014;515:587-590.

Deming DJ. The growing importance of social skills in the labor market. *Quarterly Journal of Economics* 2017;132:1593-1640.

Fischhoff B. Hindsight $\neq$ foresight: the effect of outcome knowledge on judgment under uncertainty. *Journal of Experimental Psychology: Human Perception and Performance* 1975;1:288-299.

Fischhoff B, Beyth R. "I knew it would happen": remembered probabilities of once-future things. *Organizational Behavior and Human Performance* 1975;13:1-16.

Fischhoff B. Fifty years of hindsight bias research—reflection on Fischhoff (1975). 2025.

Camerer C, Loewenstein G, Weber M. The curse of knowledge in economic settings: an experimental analysis. *Journal of Political Economy* 1989;97:1232-1254.

Paul LA. Transformative Experience. Oxford: Oxford University Press, 2014.

Kornell N, Hays MJ, Bjork RA. Unsuccessful retrieval attempts enhance subsequent learning. *Journal of Experimental Psychology: Learning, Memory, and Cognition* 2009;35:989-998.

Bjork RA. Memory and metamemory considerations in the training of human beings. In: Metcalfe J, Shimamura A, eds. *Metacognition*. Cambridge: MIT Press, 1994:185-205.

1EdTech Consortium. Open Badges Specification v3.0 (Final Release, Document Version 1.4.5, June 29, 2026); AchievementCredential JSON Schema (context 3.0.3).

Francis T Jr. On the doctrine of original antigenic sin. *Proceedings of the American Philosophical Society* 1960;104:572-578.

Gassner G. Beiträge zur physiologischen Charakteristik sommer- und winterannueller Gewächse. *Zeitschrift für Botanik* 1918;10:417-480.

Gostic KM, Ambrose M, Worobey M, Lloyd-Smith JO. Potent protection against H5N1 and H7N9 influenza via childhood hemagglutinin imprinting. *Science* 2016;354:722-726.

Hróbjartsson A, Thomsen ASS, Emanuelsson F, et al. Observer bias in randomised clinical trials with binary outcomes: systematic review of trials with both blinded and non-blinded outcome assessors. *BMJ* 2012;344:e1119.

Hróbjartsson A, Thomsen ASS, Emanuelsson F, et al. Observer bias in randomized clinical trials with measurement scale outcomes. *CMAJ* 2013;185:E201-E211.

Koedinger KR, Aleven V. Exploring the assistance dilemma in experiments with Cognitive Tutors. *Educational Psychology Review* 2007;19:239-264.

Kosmyna N, et al. Your brain on ChatGPT: accumulation of cognitive debt when using an AI assistant for essay writing. Preprint, 2025.

Leavitt JD, Christenfeld NJS. Story spoilers don't spoil stories. *Psychological Science* 2011;22:1152-1154.

Medical Research Council. Streptomycin treatment of pulmonary tuberculosis. *BMJ* 1948;2:769-782.

Moustgaard H, Clayton GL, Jones HE, et al. Impact of blinding on estimated treatment effects in randomised clinical trials: meta-epidemiological study (MetaBLIND). *BMJ* 2020;368:l6802.

Noy S, Zhang W. Experimental evidence on the productivity effects of generative artificial intelligence. *Science* 2023;381:187-192.

Rosenthal R, Jacobson L. *Pygmalion in the Classroom*. New York: Holt, Rinehart and Winston, 1968.

Salazar J, Moustgaard H, Bracchiglione J, Hróbjartsson A. Empirical evidence of observer bias in randomized clinical trials: updated and expanded analysis. *Journal of Clinical Epidemiology* 2025.

Schulz KF, Chalmers I, Hayes RJ, Altman DG. Empirical evidence of bias: dimensions of methodological quality associated with estimates of treatment effects in controlled trials. *JAMA* 1995;273:408-412.

Shannon CE. Communication theory of secrecy systems. *Bell System Technical Journal* 1949;28:656-715.

Sheldon CC, Hills MJ, Lister C, et al. Resetting of FLOWERING LOCUS C expression after epigenetic repression by vernalization. *PNAS* 2008;105:2214-2219.

Sinha T, Kapur M. When problem solving followed by instruction works: evidence for productive failure. *Review of Educational Research* 2021;91:761-798.

Slamecka NJ, Graf P. The generation effect: delineation of a phenomenon. *Journal of Experimental Psychology: Human Learning and Memory* 1978;4:592-604.

Spence M. Job market signaling. *Quarterly Journal of Economics* 1973;87:355-374.

Sympson JB, Hetter RD. Controlling item-exposure rates in computerized adaptive testing. *Proceedings of the 27th Annual Meeting of the Military Testing Association*, 1985:973-977.

The Effortless Trap: productive struggle, AI, and the illusion of learning. *arXiv:2606.26181*, 2026.

World Economic Forum. *The Future of Jobs Report 2025*. Geneva: WEF, 2025.

柏拉图. 美诺篇. 见: 柏拉图全集. 王晓朝译. 北京: 人民出版社, 2002.

论语·述而. 见: 论语译注. 杨伯峻译注. 北京: 中华书局, 2009.

同批稿: 现配优势; 零判份; 脱稿步 (V3) ; 常席差; 回改; 敞问. 2026-09-02.

第六章 单遍段

【章首】本章给方程添的第二条本体性质：**单遍性**——存在这样的工序段，尝试本身耗散其重来的条件，第二次顶着同一个名字，是另一件事。读数有二：重走落差 Δ （重做样本上二次与首次成功率之差）与门位 k （工序中首个不可撤销步的位次）。它与《盲遍》是本体层内最近的一对：盲遍管认知态的单向门（答案不能先到），本章管对象态的单向门（过程不能重来）——答案先到而对象未动的纯判断题只触发盲遍，对象已改而无外部答案的操作题只触发本章，两门各自开合。本章三场预注册检验的处置在十章里最见纪律：一场按字面可宣布过关却自愿降格为未获裁决，一场按判负条款字面判负照登并把判据设计缺陷记账。本章欠的账：跨域重做样本库（ Δ 审计）与防零膨胀的任务级判据重跑，均未做。

大白话：有些事，走过一次就回不去了

打一个鸡蛋，做舒芙蕾。它塌了。你可以再打一个新鸡蛋重做，但塌掉的这一个回不去了——不是你手艺不行，是这件事本身不接受“重来”。

生活里这类事比我们以为的多。第一次见客户的印象、第一次向父母坦白、把一条消息发出去的那一刻、手术刀切下去的第一刀、一段代码上线之后跑出去的那一笔数据。它们的共同点是：**尝试本身，消耗掉了重来的条件。**

《单遍段》讲的就是这类工序段。它的说法是：存在这样一些环节，第二次顶着同一个名字做，其实是另一件事。

为什么这在今天要紧？

因为 AI 让“重来”变得几乎免费。写不好？再生成一版。方案被否？换个角度再来一版。这本身是好事，效率翻了几倍。但它带来一个副作用：**当一切都可以重投，“通过”就不再证明什么。**

想想考驾照。如果一个人可以无限次重考，那么“考过了”这个记录说明的只是他考了足够多次，不是他会开车。以前重考是有成本的，所以考过还能说明点什么；现在在很多领域，重投的成本归零了，于是“最终通过”这个证据的价值也跟着塌了。

于是真正稀缺的，变成了**能在第一遍就走过去的的能力**——书里叫单遍力。

它不是什么？

它不是“不许试错”。恰恰相反，试错在可以重来的地方是最好的学习方式。它说的是：**你要知道你手上这段工序，哪一步是不可撤销的。**

一个厨师如果分不清“这锅汤咸了可以兑水”和“这块肉煎老了就废了”，他会在错误的地方紧张，在该紧张的地方松懈。会做菜的人心里有一张图：哪里可以回头，哪里不能。**这张图本身就是本事。**

读数是什么？

两个。一个叫**重走落差 Δ** ：同一个人、同一类任务，第二次做的成功率减去第一次做的成功率。如果一段工序真是“单遍”的，那么第二次不会明显更好——因为对象已经变了，你面对的不是同一个东西了。另一个叫**门位 k** ：这段工序里，第一个不可撤销的步骤排在第几位。 k 很小的工序，意味着你几乎一上手就没有退路了。

一个小场景。

两个学生做同一个课程项目。甲一上来就搭框架、跑通、然后不断修补；乙先做了三天调研，第四天才动第一行。中途框架要改，甲发现自己早期的一个选择已经渗透进所有模块，改不动了——那个选择就是他的门位 k ，他在第二步就走进了不可撤销区，而他当时不知道。

这不是“计划的重要性”这种老生常谈。老话讲的是要多想再动手；这里讲的是一件更具体的事：**看出哪一步是门**。有些项目就该马上动手，因为它的门在很后面；有些项目动手前必须想清楚，因为它的门在第二步。判断门在哪里，比“三思而后行”有用得多。

再看两个身边的例子。

道歉。一件事做错了，第一次道歉是有分量的。如果第一次道歉说得含糊、敷衍、带着辩解，那么第二次再郑重道歉，效果完全不同——不是因为第二次说得不好，是因为对方已经知道你会道歉了。**你消耗掉的不是机会，是对方的“未知”。**

装修。墙拆了可以补，但承重墙一旦动了，整栋楼的账都得重算。所有装修师傅都知道哪堵墙不能碰——这个知识不写在任何一本书里，它就是这一章说的“门位”。一个分不清承重墙和隔断墙的人，会在无所谓的地方犹豫，在要命的地方大胆。

三种常见的误读。

第一种：以为它在说“要谨慎”。不是。它说的是要分区：在可撤销区大胆试、快速试、多试；在不可撤销区慢下来。一律谨慎的人和一律莽撞的人，犯的是同一个错——都没有那张图。

第二种：以为重来一次就是同一件事。这是最普遍的误解，也是 AI 时代最贵的一个。你重投一版方案，看起来是“同一个任务的第二次尝试”，但你面对的对象已经变了：客户已经听过一版了，你自己的思路已经被第一版锚住了，时间也少了。**第二次从来不是第一次的重演，它是一个新的、条件更差或更好的、不同的局。**

第三种：把模拟当成单遍。演练、沙盘、模拟机、无风险的试运行，都是极好的训练手段，但它们有重来按钮。书里把这一条写得很硬：**有重来按钮的地方，不产生单遍记录。**混淆这两者的人，会以为自己练过了，直到真的没有退路那一天。

给自己做个小检查。拿出你最近做完的一件事，倒着数：第一个“做了就回不去”的步骤，排在第几位？如果你答不上来，说明那件事你其实是闭着眼睛走过来的——这次运气好而已。

它和旁边几个概念的区别。

和《盲遍》的分界上一章已经讲过：一个管认知的门，一个管对象的门，四种组合都存在。

和《峰兑》的区别在于**是不是稀发**：单遍段说的是很多平常工序里就有的不可撤销步（发一条消息、切第一刀、定一个框架），它天天在发生；峰兑说的是罕见的大事件。一次日常的不可撤销步不是峰时事件，一次峰时事件里通常包含很多不可撤销步。

和《入账之前的技能》是一对镜像：两者都是单向门，但一个在**对象侧**（做了就回不去），一个在**记录侧**（写下就归了类）。你可以私下重试很多次而始终没有入账，也可以只做一次就被永久记入某本账。

一句话记住它：这一章管的是**回不回得去**——以及你知不知道自己现在站在门的哪一边。

怎么长出来：方法、技术与流程

个人层面。

拿你手上正在做的一件事，画一条工序线，在每一步旁边标一个字：“可”或“不可”——可撤销、不可撤销。第一次做这件事的人，多半标不准；标错的地方，就是下一次要留神的地方。做过三五轮，你就有了那张图。

大学发生学教育：把“第一遍”当作教学对象。

这一条最容易被现行教育体系忽略，因为教育天然鼓励重做——重考、补交、修改后再提交，都是善意的设计。善意没错，但它使学生毕业时对“第一遍”完全没有经验。

- **首过档案。**在实践类课程里区分两种作业：**可重投作业**（改到好为止，用于学习）与**首过作业**（一次提交，不得重交，用于记录）。首过作业不必占很高分值——关键不是惩罚，是留下一个**真实的第一遍记录**。学生毕业时应该能看到自己有多少件事是一次做成的。
- **门位标注训练。**任何实践作业，要求学生先交一张工序图，标出哪几步不可撤销。教师批改这张图，比批改成品更有价值——**标错门位的地方，正是他将来会栽的地方**。
- **不可撤销实训的安全设计。**真正不可撤销的事（临床、工程、上线）必须先在有保护的环境里练，但要明确告诉学生：模拟机是“有重来按钮”的，因此它训练的是操作，不是单遍力。两者不可混为一谈——这一条书里写得很硬：**演练不构成单遍记录**。
- **失败保护条款。**首过作业失败不扣品行分、不进推荐信、不影响评奖。没有这一条，学生会用尽一切办法把首过变成预演，制度当场自毁。

教师的动作。

在课堂上示范“我现在要做一个不可撤销的决定”——比如当场选定一个方案并宣布不再回头，然后带着这个约束往下走。学生极少见过有人在他们面前这样做。

最容易做坏的地方。

把首过作业变成“一考定终身”。它的意义在于**记录**，不在于筛选；一旦用于筛选，学生会把全部精力花在把首过变成第二遍上——那正是它要测的东西的反面。

原文

单遍段：重走结构、信号坍塌与 AI 时代大学毕业生的核心竞争力

The Single-Pass Segment: Retry Structure, Signal Collapse, and the Core Competitiveness of College Graduates in the Age of AI

2026年9月2日

摘要：关于“AI 时代大学毕业生的核心竞争力将是什么”，既有回答几乎全部用存量名词作答，并共享一个从不言明的地基：通往任何存量的路都是可重走的。本文从三个互不相识的学科——恢复生态学、目击证人记忆研究、救济法——收集到同一份实证证词：存在这样的工序段，走过一次之后，“再走一次”在严格意义上不存在，存在的只是顶着同一个名字的另一件事。本文将其概念化为**单遍段**（尝试本身耗散其自身重来条件的工序段），给出三类发生机制（物理—生理改写、认知—关系覆写、制度封门）与两个可计量读数（**重走落差 Δ** 与**门位 k** ），并论证它是任务经济学的任务属性清单中缺失的一项。据此提出四个命题：可重置性是机器学习范式的可及性条件（P1）；零边际成本重投使“最终通过”的证据

值坍缩 (P2)；生成式 AI 压平可重段信号的成本差，使信号分离均衡向单遍段迁移 (P3)；由于不可逆性是制度品，评价体系将通过重装单遍环节恢复筛选功能 (P4)。本文的回答由此给出：核心竞争力不是任何存量，而是**单遍力**——在门位靠前、重走落差为负的工序上完成首过的能力，与一份无法由重投伪造的单遍履历；并指出高等教育“重试健身房”式养成管线与此的结构性错配。文末以三场预注册检验检查劳动市场边：能力暴露侧因体力性缠绕自愿降格为未获裁决；真实使用侧在申报局限内获得稳健支持；任务级词典检验按预注册判负条款字面判负、照登，并将设计教训记账。全文附带写死日期的证伪条款。

关键词：单遍段；重走落差；门位；任务框架；信号理论；恢复债；再固化；救济规则；生成式人工智能；就业能力

Abstract: Existing answers to “what will be the core competitiveness of college graduates in the age of AI” almost invariably answer with stock nouns, sharing an unexamined premise: every path to every stock can be walked again. Drawing convergent empirical testimony from three mutually unacquainted disciplines—restoration ecology, eyewitness memory research, and the law of remedies—this paper conceptualizes the **single-pass segment**: the segment of a task sequence whose very attempt dissipates the conditions of its own repetition. It identifies three generating mechanisms (physical-physiological rewriting, cognitive-relational overwriting, institutional gating) and two measurable readings (the **retry gap Δ** and the **gate position k**), arguing that retry structure is a missing attribute in the task-based economics of AI and labor. Four propositions follow: resettability conditions machine-learning trainability (P1); zero-marginal-cost regeneration collapses the evidential value of “passing” (P2); generative AI flattens cost differentials on retryable signals, migrating separating equilibria toward single-pass segments (P3); and since irreversibility is often an institutional artifact, evaluation systems will reinstall single-pass gates (P4). Core competitiveness is therefore not a stock but **single-pass capacity**. Three preregistered tests are reported in full, including one voluntarily downgraded, one supported within declared limits, and one adjudicated against the hypothesis by the letter of its own preregistration. Dated falsification clauses are provided.

JEL 分类：J24 ; J23 ; O33 ; I23

一、引言

1.1 问题与既有答案的共同地基

“AI 时代，大学毕业生的核心竞争力将是什么？”——这道题已经被回答过太多次，多到答案本身构成了一个可供研究的文体。世界经济论坛每两年发布一张“未来技能榜”，分析性思维、创造力、韧性、好奇心、终身学习轮流坐庄；咨询公司的版本把“人机协作能力”“提示词工程”排进前三；教育学的版本说是元认知与学习迁移；人文学科版本说是提问的能力、审美与同理心。这些答案彼此争吵，却站在同一块从不被检查的地基上：**它们都用存量名词作答**。竞争力被想象成一种可以持有的东西，像账户里的余额，问题只剩下“存哪种币”。

值得多看一眼的是，“核心竞争力”这个词本身的账本史就带着存量的胎记。Prahalad 与 Hamel (1990) 造出 core competence 时说的是**企业**——本田的发动机、佳能的光学，那种沉淀在组织里、对手难以模仿的集体学习。这个词后来被移植到个人身上，移植时没人检查一件事：企业的核心能力之所以成立，恰恰因为企业近似永生、可以在几十年里反复投资同一条技术轨道——它是一个**建立在无限次重来之上**的概念。把它搬到一个只活一次、每个当口只经过一次的人身上，存量思维就带着它的隐含前提一起上了船。

存量名词背后那个更深的共同假定，深到从不需要说出口：**通往任何存量的路都是可以重走的**。能力不够可以补课，证书没有可以再考，作品不行可以重做，方向错了可以转行。“终身学习”把这个假定升格成意识形态：人生被想象成一条可以无限次重新出发的跑道。大学的制度设计忠实执行着它——重修、补考、成绩替换、无限次提交的作业系统；就业市场的话语也执行着它——“先就业再择业”“三十五岁转码”。整个养成体系是一座重试健身房，它全部器械的说明书上印着同一句话：第二次与第一次是同一件事，只是又贵了一点。

生成式 AI 恰好在这个时刻到来，带来显露层的普遍拉平：谁都能在几分钟内交出结构完整、措辞得体、代码能跑的成品。于是榜单们忙着追问“机器做不了什么”，答案仍是存量名词。本文走另一条路：不问机器做不了什么东西，而问一个结构问题——在“重来”这件事上，人和机器的处境有什么根本不同；当重来的成本被打到接近零，价值会迁往哪里。

1.2 本文的主张与贡献

本文的中心主张一句话可述：**AI 对工作的接管按任务的重走结构排序，而非按脑力/体力或技能层级排序；因此毕业生的核心竞争力将不落在任何存量上，而落在“单遍段”上——每件事的工序里那段尝试即耗散自身重来条件的环节。**

具体贡献四点。**其一，概念与测量：**提出“单遍段”概念，从三个学科的实证传统中提炼三类发生机制，并给出两个可在公开记录上计量的读数——重走落差 Δ 与门位 k （第三、四节）。**其二，理论定位：**论证“重走结构”是任务经济学（Autor, Levy & Murnane, 2003 ; Acemoglu & Restrepo, 2019）的任务属性清单中缺失的一项，并以信号理论（Spence, 1973）为微观基础，给出“信号分离均衡向单遍段迁移”的机制（第二、五节）。**其三，制度预言：**由救济法传统导出“门是制度品”的命题，预言并解释评价体系的“重新装门”运动——口试与现场考核的回潮（第六节），及其与高等教育“重试健身房”养成管线的错配（第七节）。**其四，实证纪律：**以三场先于数据写死并作哈希存证的预注册检验检查劳动市场边，全文如实报告其中一场自愿降格、一场在申报局限内稳健、一场按预注册字面判负的全部结果（第八节），并给出带日期的证伪条款（第十节）。

1.3 结构安排

第二节做文献定位；第三节陈列三个学科的证词；第四节给出概念、机制类型学与读数；第五节以四个命题建立市场判断；第六节处理门的政治经济学并给出正面回答；第七节转向养成侧；第八节报告实证；第九节与近邻划界；第十节讨论局限、研究议程与证伪条款；第十一节结语。

二、文献定位：任务属性清单里缺的一项

2.1 存量与信号：竞争力研究的两条老路

关于个人在劳动市场上凭什么立足，经济学给过两个经典回答。人力资本理论（Becker, 1964）说凭**存量**：教育与经验是投资，沉淀为生产率。信号理论（Spence,

1973) 说凭**信号**：即便教育不改变生产率，只要完成教育的成本对高能力者更低，文凭就能在均衡中把人分开。半个世纪的就业能力研究基本在这两条路的延长线上加铺装。值得指出的是，两条路共享同一个未言明的前提：**存量的积累路径与信号的发送路径都是可重走的**——补考、重修、再培训、二战三战，人力资本论视之为再投资，信号论视之为再发送。没有哪一支追问过：如果存在原理上不可重走的环节，竞争力的结构会怎样改变？这正是本文的切口。

2.2 任务框架的 AI 劳动经济学：清单上没有”重走结构”

过去二十年，AI 与劳动的经济学定型为**任务框架**：职业被拆成任务束，技术按任务属性替换或补足人 (Autor, Levy & Murnane, 2003 ; Acemoglu & Restrepo, 2019) 。这个传统积累了一张著名的任务属性清单：**程序性/非程序性** (ALM 的原始轴)、**可编码/默会** (Autor, 2015 借 Polanyi, 1966 的”我们知道的比我们能说的多”解释自动化边界)、**可用监督学习/不可** (Brynjolfsson & Mitchell, 2017 的 SML 量规)、**可预测/不可预测的物理与认知工作** (Frey & Osborne, 2017 的自动化概率)。测量端则有 Felten, Raj & Seamans (2021) 的能力重叠指数与 Eloundou 等 (2023) 的任务暴露评级。

把这张清单排开看，缺口就显出来了：**清单上没有任何一项刻画任务的重走结构**——尝试失败之后，这个任务还是不是原来那个任务。默会性问”能不能说清”，可编码性问”能不能写成规则”，SML 量规问”有没有可学的输入输出对”，它们全都默认任务对象在尝试之间保持恒同、试错可以清场重来。而这个默认恰恰在生成式 AI 时代变成了承重墙：本文第五节将论证，重走结构不是清单上可有可无的一列，而是决定这一代技术接管顺序的那一列——因为这一代技术的训练范式 (可重置环境) 与部署优势 (零边际重投) 都长在重走结构上。

2.3 不可逆性诸传统：各占何处，空隙何在

“不可逆”当然不是无主之地。环境经济学有准期权价值 (Arrow & Fisher, 1974) 与不可逆投资理论 (Dixit & Pindyck, 1994)：当行动后果不可逆而信息会到来，等待本身有价值。决策哲学有 Shackle (1949; 1961) 的”关键实验”：有些决定”实验行为本身可能永远摧毁其进行时的环境”，因此概率论对它们失语。社会学有 Goffman (1967) 的命运性时刻，政治哲学有 Arendt (1958) 的行动不可逆性与作为救济的宽恕、许诺。晚近则有直接以 AI 为题的”后果不可逆性”论 (Metis AI, arXiv:2605.14407, 2026)：后果会不可逆传播的任务抗自动化。

这些传统合起来占住了”不可逆”的哲学表述与**后果轴**——行动的效果在世界里收不收回。它们共同留下的空隙有三：其一，无人区分后果的不可逆与**任务对象自身被尝试改写** (第4.5节将证明这是两根正交的轴)；其二，无人给出可在公开记录上计量的读数；其三，无人把这根轴接到劳动分工与养成管线上。本文在这三处落脚。

三、三个学科的证词

理论定位讲完，先听证。以下三个学科互不引用、数据互不相通，却各自用自己的实证材料撞上了同一堵墙。

3.1 恢复生态学：第二次走到的是另一个地方

恢复生态学是一门以”第二次尝试”为全部内容的学科：生态系统被破坏了，人类回头重走一遍，看能不能走回去。Bradshaw (1987) 称之为生态学的”酸性试验”——你说你懂一个系统是怎么长成的，那你把它重新长一遍试试。学科的行业标准 (国际生态恢复学会

SER 各版《生态恢复原则与标准》，Gann et al., 2019) 围绕“参照生态系统”组织作业：先确定“它本来的样子”，再把受损系统往参照上修。这套记法的底层假定与大学的重修制度一模一样：路是可重走的，第二次是第一次的复刻，差别只在成本与保真度。

然后是这门学科自己交出的数据。Moreno-Mateos 等 (2012) 对全球621个恢复与新建湿地做元分析：即使在恢复开始后五十到一百年，湿地的生物结构仍比参照低约26%，生物地球化学功能低约23%。Benayas 等 (2009) 覆盖更广生态系统类型的元分析给出同一形状：恢复显著提升了多样性与生态系统服务，但仍系统性低于未受损参照。更要命的是“还差一截”，而是恢复曲线的**形状**：不是缓慢逼近参照的渐近线，而是**收敛向另一个地方**。Suding 等 (2004) 的“替代稳态”研究表明，被推离原路径的系统会被正反馈锁进新状态；Fukami (2015) 的优先效应研究把机制推到最尖锐处：物种**到达的次序**决定群落装配走向——次序不可重排，所以装配不可重演。这门学科给这笔账起了名字：**恢复债**

(Moreno-Mateos et al., 2017) ——不是还没还清的债，是**结构上还不清的债**。学科内部“新生态系统”派 (Hobbs et al., 2009) 与参照保真派 (Murcia et al., 2014) 的路线斗争，吵的正是要不要停止假装第二次是第一次的复制品。

证词：**第二次不是慢一点的第一次，是另一件事。**

3.2 目击证人记忆研究：取回即改写

这门学科早八十年撞上同一堵墙，撞的对象是人自己的过去。Münsterberg (1908) 已警告法庭：记忆不是档案柜，提取不是调卷。Loftus 与 Palmer (1974) 的实验里，仅把提问动词从“碰撞”换成“猛撞”，证人对车速的记忆就系统性上移，一周后还“记得”并不存在的碎玻璃；Loftus、Miller 与 Burns (1978) 的误导信息范式给出机制：事后信息写进原始记忆，且写入之后**无法分离**——“证人真诚地”记得“他没见过的东西。神经科学补上底层：再固化研究 (Nader, Schafe & LeDoux, 2000) 表明记忆每被提取一次，痕迹就进入可改写状态再重新固化——**回忆不是读取，是覆写**。第二次回忆读到的，是被第一次回忆改过的版本。

司法应用把“单遍”直接写进程序：证人辨认原则上只做一次——重复辨认会系统性抬高错误指认率，因为第一次辨认已把“我在队列里见过这张脸”写进证人记忆，第二次辨认辨的不再是案发现场，而是上一次辨认 (Deffenbacher et al., 2006 的元分析；美国司法部1999年《目击证据指南》以来的单次辨认规则)。程序的回应不是“多试几次取平均”——平均只会把污染平均进去——而是承认这件事只能做一遍，把全部工艺压进这一遍里。

还有一面镜子直照恢复生态学：上世纪的“恢复记忆治疗”相信被压抑的记忆可以像湿地一样修复回参照状态，结局 (Loftus, 1993 及“商场走失”系列实验) 是修复程序制造出崭新的、当事人深信不疑的假记忆，无辜者因此入狱。**以参照为目标的第二次尝试，产出的是新东西，还坚信自己是原物**——与替代稳态是同一句话在两个学科的两度发音。

3.3 救济法：能不能重来，是规则场指定的

前两家说的是事实层：第二次到不了原处。法学把问题整个倒过来：社会里大量的“不可重来”根本不是自然事实，而是**制度制品**。Holmes (1897) 的冷话揭开这一层：合同不过是“要么履行、要么赔钱”的选择权——法律可以把任何事情**指定**为可重来 (赔钱了事) 或不可重来 (覆水难收)。Calabresi 与 Melamed (1972) 把指定方式列成三格：**财产规则** (未经同意不得动)、**责任规则** (可以动，按官定价格赔——即一切皆可重来，只是有价格)、**不可让渡规则** (给多少钱都不许动)。一件东西落在哪一格，不是它的物理属性，而是立法与裁判的选择。

这套以“万物可定价”为气质的框架自己留了两个漏洞，泄露出规则指定不了的东西。其一是特定履行之争（Kronman, 1978；Schwartz, 1979）：什么时候赔钱不够、必须强制履行原合同？传统答案是“标的物独一无二”——法律在此被迫承认有些东西责任规则买不回来。Radin（1987）的“市场不可让渡性”把第三格推到底：有些东西一旦允许定价，其性质即告改变——**定价这个动作本身就是一次不可逆的改写**。其二是复职救济的实证：不当解雇者被判复职后，相当比例在一两年内自行离职——法律有权命令第二次，但命令不出“同一件事”。

法学与生态学甚至已在现实里正面相撞过一次：**湿地缓解银行**。美国《清洁水法》体系允许开发商填掉一块湿地，只要购买“缓解信用”——由专业机构在别处修复或新建等量湿地。这是教科书级的责任规则操作：生态系统被记成可互换的英亩数。而 Moreno-Mateos 的621块湿地正是对这本账的审计：**信用簿上两清了，恢复债还在**。法律的记法宣布重来完成，生态的数据显示重来从未发生——同一块湿地，两本账。

三个学科，三套互不相通的数据，同一句话：**存在这样的工序段——走过一次之后，“再走一次”在严格意义上不存在；存在的只是另一件事，顶着同一个名字**。

四、单遍段：概念、机制类型学与读数

4.1 定义

定义1（单遍段）：一件事的工序序列中，凡**尝试本身耗散其自身重来条件**的工序段。第一遍执行改写了执行对象，“第二遍”接手的已是另一个对象——无论记录多么完整、时间多么充裕、预算多么慷慨。其反面为**可重段**：尝试不改写对象，失败可清场重来，第二遍与第一遍是同一件事的两个实例。

日常生活里两类段落随处可见。写简历是可重段——改一百遍，纸不记仇；**面试的第一分钟是单遍段**——开场已写进考官的知觉，第二次“重新自我介绍”面对的是被你的第一次改变过的考官。改论文是可重段；**答辩被问倒的那一刻是单遍段**——可以补充陈述，但“当场答不上来”已经发生，补充是在这个新事实之上进行的另一件事。谈判可排练一百次，全是可重段；**开出的第一个报价是单遍段**——锚一旦抛下，本局水文即改，只能在被它改过的水里继续游。急诊医学把这件事量化到残酷：气管插管首次尝试失败后，气道水肿、出血、视野恶化，每多一次尝试并发症发生率陡增（Sakles et al., 2013），这个行当因此把“首过成功率”当核心质量指标——**不是因为医生第二次会手抖，而是因为第二次插的已不是第一次那条气道**。

注意定义的颗粒度：单遍性不是任务的形容词，是**工序段**的属性。几乎没有哪件事整个儿是单遍的——谈判是无穷次排练（可重）加一次开价（单遍），手术是模拟器上的一千次练习（可重）加真皮肤上的一刀（单遍）。每件事都是两种段落的编织物。这个颗粒度是第五节全部市场判断的前提：AI 不按“职业”接管，按段接管。

4.2 机制类型学：单遍性从哪里来

三家学科恰好各贡献一类发生机制，合成表1。

表1 单遍段的三类发生机制

机制	载体	原型案例	学科来源
	任务对象的物质状态		

M1 物理—生理改写		首次插管后的气道；被扰动湿地的土壤化学与种子库；发掘即毁的地层	恢复生态学
M2 认知—关系覆写	交互对方（或本人）的知识与关系状态	被提取即改写的记忆；抛出即锚定的首个报价；说破即回不去的表白；披露即不可收回的信息	目击证人记忆研究
M3 制度封门	规则场的登记状态	既判力与一事不再理；临床试验揭盲；V1 之后必须起飞；应届生身份窗口	救济法

M2 值得单独强调一个纯逻辑的骨架：**知识状态的单调性**。对方（或你自己）一旦知道了某件事，就无法退回不知道——信息可以被否认、被遗忘表层，但”已被告知”这个事实性改写不可撤销。这解释了为什么人际工序密布单遍段：凡是以改变对方知识或关系状态为内容的动作（报价、坦白、威胁、道歉、第一印象），执行即改写交互对象。M3 则提醒：三类机制常常复合——应届生身份（M3）之所以焦灼，是因为它叠在”第一份工作塑造职业轨道”（M1/M2）之上；反过来，M3 也可以被单独装拆，这是第六节的支点。

4.3 读数一：重走落差 Δ

定义2（重走落差）： Δ = 二次尝试的成功率 - 首次尝试的成功率，在同类事件的重做记录上统计。

Δ 的符号把世界切成两半。 **$\Delta \geq 0$ 的世界**：申请与投稿——NIH 历年统计里按评审意见修订后重投（A1）的资助率显著高于首投（A0）；代码——环境每次重置如新，此即”迭代”的本义；标准化考试——重测效应元分析显示重考者平均提分（Hausknecht et al., 2007）。这半边世界重来不但可能，而且越重来越好：对象不记仇而记录在积累。 **$\Delta < 0$ 的世界**：插管、初诊首次接触、证人辨认、危机首响应、现场演出、首个报价、第一印象。这半边对象记仇：每次尝试都在消耗下一次的条件。

计量上须交代一处方法学暗礁：重做样本天然带**选择偏差**——走到第二次的都是第一次没成的，而第一次没成与难度相关。故 Δ 的估计要么在可比难度层内做（同一医院同一适应证的插管记录、同一机构同一学科的申请档案），要么明示它测的是”实际发生的二次尝试”的处境而非反事实。第十节的清点手册把这条口径写死。

4.4 读数二：门位 k

定义3（门位）： k = 一件事的工序序列中，第一个不可撤销步的位次。

原型来自把单遍性管理得最成熟的行业：每次起飞前机组计算 V1——决断速度，滑跑中一旦超过，中断比继续更危险，门在此关闭。V1 不是形而上学，是当日风速、温度、跑道、载重代入性能手册算出的数，写进检查单、滑跑时大声喊出。 k 越靠后，事情越”可草稿”：写作的门位极靠后（发表前一切可改）；临床试验的门位在揭盲；诉讼的门位在既判力发生处。 $k=1$ 的事从第一步就玩真的：开口即锚的谈判、下针即创的操作、开场即定调的课堂。

门位一经定义，立刻显出被系统性忽视的事实：**门的位置常常是制度装的**。签署使文件定稿——签署制度发明之前”定稿”并不存在；宣誓使证词与闲聊从此两个量级；揭盲程序是统计学家发明的人造门；”一事不再理”是为判决终局性拧上的螺丝。反之，版本控制、撤回键、无限草稿文化、成绩替换政策是把门往后拆。门可装可拆——这是第六节的全部支

点，也是三家里法学那一维的独有贡献：前两家只能说门在哪儿，法学说门是谁装的、还能装在哪儿。

4.5 两根正交的轴

单遍性（对象自耗）必须与既有文献的”后果不可逆”轴严格区分。二者正交，可交叉出表2的四格。

表2 后果不可逆 × 对象自耗：四格与例证

	任务对象可重走 ($\Delta \geq 0$)	任务对象自耗 ($\Delta < 0$)
后果可逆/轻	打草稿、改代码、做习题	谈判首个报价；社交首因印象
后果不可逆/重	法官判案（后果极重，但明天还有新案、每案对象崭新，手艺在千百次重复中打磨）	首次插管；危机首响应；V1后的起飞

左下格是关键证据：**后果沉重而任务可重走**的工作大量存在，判案、签发处方、批贷款皆是——若”抗自动化”只看后果轴（Metis AI 的立场），无法解释为什么这些工作的文书与分析环节正被迅速接管而其单遍环节（庭审、面签、现场尽调）没有。右上格同样是证据：**后果轻微而对象自耗**的段落（首个报价）机器至今不能代走首过。可见接管的排序变量是列（ Δ ），不是行（后果）。

4.6 一条必须先钉死的界线：贵 ≠ 单遍

烧掉一栋楼损失惨重，但保险赔付后可原址重建同一张图纸——这是贵，不是单遍。判别法一句话：**出钱买不买得回”同一件事”**。凡钱能买回的，都在责任规则的格子里，属 $\Delta \geq 0$ 半边，无论标价多高；单遍段的定义性特征是不在任何价格上出售第二遍——621个湿地上钱花了一百年，买回来的是另一个生态系统。保险业无意间为界线提供了产业级确认：可保的前提是风险可分列、可入大数——正是 Shackle 的”可分列实验”，同类事件颇多，单次损失在池子里摊平，损失因此在**金钱维度上可重来**；保险业拒保或只按金钱等价物赔付的东西——独一无二之物、名誉的第一次破损——恰落在单遍半边。把”高风险””高成本””高利害”混同于单遍，是这个概念最容易死掉的方式，本文先于一切批评者把它钉死。

五、四个命题：AI 按重走结构接管

命题1（训练可及性）：一类任务能被当代机器学习范式大规模习得的前提，是它能被表示为可重置的试错环境；对象自耗的工序原理上拒绝这种表示。

论证要点：预训练吃可反复重读的静态语料；强化学习的试错以”每次尝试后环境回到原点”为存在条件——奖励可反复结算、轨迹可反复采样。一个尝试即耗散自身条件的环境无法提供这种采样：你不能在同一个证人身上试一百万种问法，第三种问法之后那个证人已经没有了。模拟器能缓解，但模拟的本质恰是把 Δ 拉回非负——模拟器的存在理由就是”可以重来”——于是在模拟中习得的策略对真实单遍段的迁移，本身构成一个未解的”模拟到现实”缺口。■

命题2（重投与证据坍缩）：当重投的边际成本趋近于零且重投次数不入记录，”最终通过”不再携带关于行为者的信息。

论证要点：统计学早已为此发明形式语言。序贯分析的 α 消耗函数（Pocock, 1977；O’Brien & Fleming, 1979）规定临床试验每提前看一眼数据都要从有限显著性预算里扣

账，因为看的次数本身消耗证据；方法学界把不记账的重试称作“研究者自由度”与“分叉小径的花园”（Gelman & Loken, 2014）——足够多的重试可让任何东西通过任何检验。生成式 AI 把这个统计学诅咒平民化：任何人可对任何交付物重投到通过为止，而交付物不显示重投次数。一遍过与五百遍过交出同一份光洁成品，成品上没有次数的指纹。■

命题3（信号分离均衡向单遍段迁移）：生成式 AI 压平可重段信号的发送成本差，使建立其上的分离均衡坍塌为混同均衡；仍能支撑分离的信号只剩单遍段上的表现，因为首过无法加购。

论证要点：Spence（1973）的机制里，信号（文凭、作品）之所以能把人分开，是因为发送成本与能力负相关——低能力者模仿信号的成本更高，故均衡中不模仿。生成式 AI 对一切可重段信号（文书、代码、设计稿、标准化答案）做了两件事：把发送成本的**绝对水平**压到接近零，把发送成本的**人际差**压到接近零——人人皆可重投到通过，成本差消失，分离条件破坏，均衡混同：雇主看到光洁成品，更新不了对人的信念。此即“显露层拉平”的微观基础。但有一处成本差是重投压不平的：**单遍段上的首过表现**。它无法加购（重走落差为负，买第二次买到的是另一件事）、无法外包给机器（命题1）、无法靠重投伪造（命题2），于是它对不同能力者保持不同的“成本”（成功概率），分离条件在此存续。信号市场的均衡因此整体向单遍段迁移——不是因为单遍段高贵，而是因为那里的信号还没被伪钞污染。■

命题4（接管排序）：由命题1—3，AI 对工作的实际接管顺序按任务段的重走结构（ Δ 的符号与门位 k ）排序，而非按脑力/体力、技能层级或后果轻重排序。

这四个命题合起来回答了“为什么以前没人需要这个概念”：在 AI 之前，可重段与单遍段的**相对价格**长期稳定，分不分家不值钱；这一代技术单方面把可重段价格砸穿地板，两类段落的价差第一次大到成为劳动市场的一级结构。概念是被相对价格的历史剧变逼出来的。

六、门的政治经济学与本文的回答

6.1 重新装门：一个可下注的制度预言

命题5（装门）：当可重段信号全面混同（命题3），任何仍需履行筛选职能的评价体系，将在其工序序列的有限门位处重建单遍环节。

这不是预言家的手势，是正在发生的清单：ChatGPT 之后，各国大学成规模恢复现场手写考试与口试；技术面试从“带回家的作业”改回监考白板与现场结对；答辩加重现场追问权重；执业认证收紧重考间隔与次数。每一项都是把门位 k 往前拧——把评价从“交上来的东西”（可无限重投，信号已死）移到“当场只发生一次的事”（重投不可能，信号还活着）上去。机制正是命题3的逆用：评价体系要恢复分离均衡，只能把信号环节搬进重投够不着的地带。

在这个坐标里重看熟悉的制度，会看见本相。**试用期**是制度装出来的第二次窗口——劳动法在默认单遍的世界里（招错人、入错行代价近乎不可逆）人工开凿的 $\Delta \geq 0$ 缓冲带；它的存在本身就是“世界默认单遍”的证据，正如救生圈的存在证明水会淹死人。**成绩宽恕**

（重修后新成绩覆盖旧成绩）是反向操作——记法出面否认第一次发生过。**应届生身份**是中文世界最熟的一扇装出来的门：校招通道、部分公务员岗位、落户指标只对“第一次走出校门”的窗口开放，错过即永久关闭；本文不评价这扇门装得对不对，只指出它证明了门的可装性——一个纯粹的行政类别硬生生造出一段全社会公认的不可重走。**口试的回潮**是最直白

的装门：考官要看的不是正确答案（可重段，已通胀殆尽），而是你在不可排练的追问序列里首过的样子。

6.2 装拆的分配政治

装门与拆门不是中性的技术选择。**拆门有利于供给方**：无限草稿、可撤回、可重修，降低每次尝试的赌注，让更多人敢进场——这是几十年教育民主化与互联网文化的共同方向，平权收益真实；但拆门同时让显露通胀，最终伤害的恰是无力购买“重投次数”之外信号的人。**装门有利于需求方**：门越靠前，一次表现携带的信息越多，筛选越便宜；但门放大运气与状态波动的暴政——一考定终身的旧暴力正是被 Bloom 们拆掉的。所以未来不会是全装或全拆，而是一场**门位的再谈判**：哪些环节必须一次定音（那里需要真信号），哪些环节保留无限重来（那里需要参与和练习）。毕业生将生活在这场再谈判的前线。

6.3 正面回答

AI 时代，大学毕业生的核心竞争力将是单遍力：在门位靠前、重走落差为负的工序段上完成首过的能力，以及一份由真实首过累积而成、无法由重投伪造的单遍履历。

它不是存量名词——“无法”拥有”后放着保值，只存在于一次次不可排练的执行里，像现场感只存在于现场；也不是某个行业——每份职业内部都有自己的单遍段与可重段，接管吃掉后者，人被重新定价的地方是前者。竞争力不迁往某类工作，而迁往每类工作里的那一段。“哪些职业是安全的”从此是提错了的问题；对的问题是：**你打算靠哪一段吃饭。**

6.4 附论：对”选什么专业”的重读

这个回答顺带重写了一个每年折磨千万家庭的问题。“AI 时代选什么专业安全”在本文框架里是提错了的问题——安全从来不属于专业，属于段；但专业之间确有两个可比的实际参数。**其一是单遍密度**：一个领域的核心工序里 $\Delta < 0$ 段占多大比重。临床医学、庭审律师、口译、外交、销售、表演、危机公关、中小学课堂的单遍密度高——它们的核心兑现环节是床边、法庭、同传箱、谈判桌、舞台、教室，全是首过现场；笔译、财务合规文书、标准化检测报告的单遍密度低——核心交付物几乎全落在可重段。**其二是配给渠道的存量**：该领域是否还保有把新人送进真实单遍现场的制度化通道。医学的规培与床边教学是现存最完整的单遍配给体系；法学的诊所教育次之；多数文科与商科专业的单遍配给几近于零——学生四年不曾在任何真实的 $\Delta < 0$ 现场首过一次，全部训练发生在重试健身房里。两个参数一乘，专业间的分化立刻清晰：**高单遍密度且配给体系完整的专业，其学位在 AI 时代的定价能力最稳；低单遍密度或配给已断的专业，学位将加速退化为混同信号**——不是因为知识没用了，而是因为该学位能证明的一切都落在了可被重投伪造的半边。对个人的操作含义不是换专业，而是无论在哪个专业里，都去找到并占住它的单遍段与配给渠道。

七、养成管线的错配

这个回答照出教育侧的结构性错配：**大学是重试健身房，而重试恰恰是机器的母语。**

先说公道话。把教育改造成可重试的，是二十世纪教育学最大的人道主义成就之一。Bloom (1968) 的掌握学习纲领——时间变量化、允许重学重考直到达标——拆掉了”一考定终身”的暴力；其背后信念（几乎所有人都能学会几乎所有东西，只要允许足够多遍）在 $\Delta \geq 0$ 的领地千真万确。本文不主张拆掉健身房，指出的是它的**覆盖面幻觉**：健身房默认

自己训练的能力可迁移到一切场合，而它训练的那半边——在可重置环境里迭代到通过——正是 AI 原生的半边。一个在无限次提交系统里长大的学生，被系统性训练去做的事，机器比他做得更快更便宜；而他将被市场重新定价的时刻—— $k=1$ 的时刻——健身房里没有一台器械对应。他练了四年深蹲，上场却是跳水。

四条推论，每条可检验、可操作。

推论一：单遍环境是配给品，且正在更稀缺。床边首次问诊、真实客户的首次提案、有真实观众的首演、动真金白银的首谈——这些环境无法用模拟完全替代（模拟的本质是把 Δ 拉回非负），历来靠学徒制配给：住院医、律所初级岗、剧团龙套、销售陪访。而 AI 正在接管的初级可重段工作，恰是新人**换取单遍环境入场券的货币**——初级律师靠整理卷宗换旁听谈判的资格。可重段被机器拿走，入场券作废，学徒梯子的下面几级正在消失。于是单遍环境的配给将成为教育机构的核心职能与稀缺资产：诊所式教学、真实委托人的法律诊所、带真实预算的项目制、有校外观众的公开演出与答辩——谁能把学生送进真实的 $\Delta < 0$ 现场并兜住后果，谁的毕业生就带着单遍履历出校门。这比任何课程改革都昂贵，也都难以被 AI 平替。

推论二：首过权需要保护。把每个问题的首过让给 AI 的学生，失去的不是知识——知识可后补，属 $\Delta \geq 0$ ——失去的是**首过本身，而首过不可补**。看过答案的人无法赎回未看过答案的自己：这是再固化研究在学习上的直接推论（接触即覆写且无法分离），也是每个做题人的体感——“瞄过解答再”独立”做出的题，和硬啃下来的题，在你身上留下两种东西，而你再没有机会对同一道题硬啃。须与流行的”AI 依赖致退化”担忧区分：退化论说余额问题——用进废退，废可再用回来，原理可逆；本文说单遍问题——每个问题只对你敞开一次未被剧透的第一遍，它是消耗品。退化可靠戒断修复，首过没有修复程序。养成的工艺问题因此变成一门注定要出现的学问——**首过配给学**：哪些问题的首过留给学生自己走（大抵是他打算靠它吃饭的主线），哪些大方交给机器（主线之外的一切支线）。挥霍首过与囤积首过同样失败：前者养出没有单遍履历的人，后者养出效率上被同代碾压的人。

推论三：中断后的养成不回原路，且不必伪装要回。间隔年归来者、二战三战的考生、三十岁转行者——重试健身房的记法把他们登记为”回到跑道”，用”落后几年”计量。恢复生态学的镜子照出另一本账：他们是替代稳态——第二次走的是另一个人在走另一条路，土壤已换、优先效应已把装配次序永久改写，渐近线不指向参照。这是记账方式的解放：对这些人的评价不该问”离原参照差多少”（恢复债结构上还不清，此路必然以羞辱收场），而该问**新稳态自身立不立得住**——新组合有没有功能、能不能自持、是否长出原路径上长不出的东西。这正是新生态系统派对参照保真派的胜利在人身上的重演。

推论四：单遍履历有记法悖论，需要证人结构。单遍力的证据天然难以书面化：一旦把首过转写成可无限复制的证书或作品，它就流回可重段的通胀区——写在纸上的东西都可被重投伪造。传统社会早有对策，只是原理被遗忘：**推荐信不是证书，是证词**——具名者以信誉作保，声称亲眼见过你的某次首过。证词的效力恰恰依赖第二家学科揭示的脆弱性：一个证人的一次记忆，会褪色、会被污染、不可批量再生产——而正因不可批量再生产，它不通胀。AI 时代的单遍履历将更依赖证人结构：具名的、在场的、以信誉连带的叙述，而非可下载的凭证。操作含义直白得近乎古老：**你的每一次首过，都要有证人。**

八、实证：三场预注册检验

8.1 设计与纪律

本文第五节的市场判断有一条最便宜的可检验边：若接管按重走结构排序，则错误后果越难纠正的职业，与生成式 AI 的交集应当越小。我做了三场检验，每场的预注册均在读取该场数据之前写死并作 SHA-256 哈希存证（预注册一：5410b75b...87916b；预注册二：780a733c...1eec35；预注册三：7c63f873...3220ab1fa），判负条款、不利结果处置与停止规则均先于数据落笔。前两场共用同一个不可逆性代理：O*NET 30.0 的“错误后果”量表（Consequence of Error——“若工人犯下**不易纠正**的错误，后果通常有多严重”，元素 4.C.3.a.1；预注册一曾笔误作 4.C.3.a.3，构念以名称唯一指认，仅更正编号），条件于教育准备层级（Job Zone）。事先申报的构念错位一并交代：该量表偏向“后果严重度”，只能间接代理“对象自耗”轴——它是公开数据里能取到的最近代理；检验的是市场判断的劳动市场边，不是承重命题本体。

这一节的存在方式本身是论点的自我演示：预注册就是给自己装门——看见数据前把判负条款签字画押，就是把分析的门位拧到 $k=1$ ，让“看一眼”从此要记账（命题2的逆用）。没有这扇自装的门，下文那些负相关可以被无限重投出来，而它们将一文不值。

8.2 第一场（能力暴露侧）：自愿降格为未获裁决

以 Felten 等（2021）的 AIOE 指数为因变量——它计量职业能力构成与 AI 技术能力的重叠。合并样本682个职业。主检验过线：条件于 Job Zone 的偏 Spearman 相关 **-0.34**（单侧 $p \approx 4 \times 10^{-20}$ ）；错误后果十分位从低到高，AIOE 均值从 +0.45 单调降至 -0.49。但预注册的唯一稳健性检验不留情面：加控体力活动重要性（O*NET 4.A.3.a.1）后，偏相关塌缩至 **-0.015**。预注册字面写“符号仍为负则稳健”，-0.015 字面仍负——我拒绝按字面宣布过关：量级塌缩的实义是“错误不可纠正”与“体力性”在当今职业结构里深度缠绕（错误最难挽回的职业——飞行、外科、急救——多半也最动手），这场检验分不开两者。第一场**自愿降格为“未获裁决”**：数据与“按重走结构排序”相容，也与“AI 只是避开体力”相容。

8.3 第二场（真实使用侧）：在申报局限内稳健支持

换掉因变量的性质：不问“AI 能力与这职业重叠多少”，改问“**人们实际拿生成式 AI 干了这职业的多少活**”——以 Anthropic Economic Index 公布的职业级实际使用暴露（基于真实对话的任务映射）为因变量， $n=668$ 。预注册二写明：若就业基数可精确匹配则用使用占比对就业占比的对数比；实际匹配失败（公开就业表按职业名而非编码给出，精确匹配率为零），按预注册回落口径改用对数使用暴露，**未除以就业基数的局限如实申报**——职业越大，出现在使用记录里的机会机械地越多。结果：主检验偏 Spearman **-0.298**（单侧 $p \approx 2 \times 10^{-15}$ ）；关键在稳健性——**加控体力后 -0.128**（单侧 $p \approx 5 \times 10^{-4}$ ），符号与量级俱存，十分位表单调无反转。能力暴露指数分不开的东西，真实使用记录分开了：**教育层级与体力性相近的职业之间，错误越难挽回的那个，人们实际让 AI 沾手越少。**

8.4 第三场（任务级词典检验）：按预注册字面判负，照登

针对“任务级证据缺席”的短板，第三场下到任务层：以先于数据写死的两部正则词典把 O*NET 任务语句分为“单遍指向”（negotiat/testif/emergenc/rescue/surger/mediat 等16词根，命中449条）与“可重指向”（prepar/draft/compil/analyz/review 等12词根，命中3,083条），匹配 Anthropic 任务级渗透数据（15,714条任务），对同时含

两类任务的229个职业计算逐职业均值差 d ，预注册判据： $\text{median}(d) \leq 0$ 且 Wilcoxon 单侧 $p \leq 0.05$ 为支持， **$\text{median}(d) \geq 0$ 判负**。

结果： $\text{median}(d) = 0.0000$ ，**判负条款字面成立，本场判负，照登**。判负的技术根源是我自己的设计缺陷：任务渗透度高度零膨胀（大量任务渗透为零），把中位数吞成了零——判据未对零膨胀设防。同批预注册统计量中的 Wilcoxon 符号秩（单侧 $p \approx 7 \times 10^{-9}$ ）与稳健性任务层比较（单遍典任务渗透均值 0.041 对可重典 0.153，Mann-Whitney 单侧 $p \approx 4 \times 10^{-12}$ ）方向强负，但按停止规则，它们只能作为描述登记，不改判决。第一场我拒绝按字面占便宜，第三场就按字面认栽——这是同一扇门的两面。

8.5 合并裁决

表3 三场预注册检验汇总

	第一场（能力暴露侧）	第二场（真实使用侧）	第三场（任务级词典）
因变量	Felten AIOE（能力重叠）	AEI 职业级实际使用暴露	AEI 任务级渗透度
样本	682 职业	668 职业	15,714 任务 / 229 职业
主检验	$\rho = -0.34$ ($p \approx 4 \times 10^{-20}$)	$\rho = -0.298$ ($p \approx 2 \times 10^{-15}$)	$\text{median}(d) = 0.0000$
稳健性（加控体力）	塌缩至 -0.015	-0.128 ($p \approx 5 \times 10^{-4}$)	（任务层均值 0.041 vs 0.153，仅描述）
预注册裁决	自愿降格：未获裁决	支持且稳健（限内）	判负（照登）
记账事项	体力缠绕不可分离	未除就业基数；单模型样本	判据未防零膨胀（设计缺陷）

按预注册二”结论以更保守者为准”的约定：**市场判断在真实使用侧获得稳健支持（在两条已申报局限内：未除就业基数、样本限单一模型用户）；在能力暴露侧未获裁决（体力缠绕未分离）；任务级首次尝试按预注册判负（判据设计缺陷记账），任务级裁决仍然开放——第十节 F3 为它保留赌约。一条标明 事后探索、不计入任何判定的补看：第一场低体力半样本内（ $n=341$ ）负关联仍在（ $\rho = -0.18$ ， $p = 0.001$ ），高体力半样本为零——仅具启发值。**

九、与近邻的分界

与”后果不可逆”论（Metis AI, 2026 ; Arrow & Fisher, 1974 ; Dixit & Pindyck, 1994）：其轴是行动后果在世界里的传播与保留选择权的期权价值；本文的轴是尝试对任务对象自身的耗散。两轴正交（表2四格）：法官判案后果极重而任务可重走，首个报价后果轻而对象已改。期权理论教人为保留选择权而等待；重走落差教人认清有些事等多久都只有一遍——等待只是把首过推迟到条件更陌生的一天。

与默会知识轴（Polanyi, 1966 ; Autor, 2015）：波兰尼悖论解释自动化边界靠”说不清”——默会的任务难编码。默会与单遍同样正交：骑自行车高度默会而完全可重走（摔了再上）；按剧本开出首个报价毫不默会而严格单遍。ALM—Autor 传统的清单缺的正是重走列，而这一代技术的接管边界（命题1、2）恰好长在这一列上。

与任务替代框架（Acemoglu & Restrepo, 2019）：该框架给出替代与新任务创造的一般均衡账，但其任务空间未刻画重走结构。本文可读作向该框架提交一个新的任务属性与一组测量（ Δ 、 k ），并预测该属性将成为替代函数的一级自变量。

与信号理论 (Spence, 1973)：非对手而是积木。本文的贡献是指出生成式 AI 系统性破坏可重段信号的分离条件（成本差压平），并推导均衡向单遍段迁移（命题3）——这是信号理论在重投成本趋零世界里的延拓。

与 Shackle 的关键决策 (1949; 1961)：他占住机制的哲学表述权（“实验行为本身可能永远摧毁其进行时的环境”），本文如实承认；接过的是他没做的三件事：可计量读数 (Δ 、 k)、接到 AI 时代的劳动分工、接到养成管线。

与遍历性经济学 (Peters, 2019)：共享”重复不可随意换序”的形式直觉，但其轴是财富动力学中的不可遍历（乘性过程与破产吸收态），救济是换效用函数；本文的轴是任务对象被尝试改写，读数落在重做样本与工序位次。

与 Goffman 的命运性时刻 (1967)：其轴是道德展演与性格证明；且未区分”高利害”与”单遍”（4.6节），而这条界线是本文市场判断的承重墙——没有它，“AI 该避开一切高风险场合”的平庸结论会冒充本文。

与 Arendt (1958)：她把行动不可逆当人的境况，把宽恕与许诺当救济。本文把她的救济接到卡拉布雷西—梅拉米德三格上，翻出她未翻的一面：救济即装门拆门的制度技术，不可逆的地形可政治操作——此即命题5的地基。她给了不可逆哲学的重量，本文给它一张可施工的图纸。

与本雅明的灵光 (1935)：机械复制使艺术品的”此时此地”性枯萎；生成式 AI 是那篇文章的续集，复制对象从艺术品扩大到一切交付物。但其轴是本真性与仪式价值的文化命运；本文计算灵光退守到哪一段、市场愿为那一段付多少钱。

与现场性之争 (Phelan, 1993 ; Auslander, 1999)：表演学争现场不可复制性的本体论与政治价值；本文不站队，只借双方公认的结构事实（现场只发生一次），给它一个双方都没给的劳动市场判据。

与刻意练习传统 (Ericsson et al., 1993)：一万小时论是重试健身房的科学宪章——刻意练习的定义性特征恰是可重复、有反馈、容许失败。Ericsson 自己严格区分练习与表演：练习在 $\Delta \geq 0$ 区制造能力，表演在 $\Delta < 0$ 区兑现能力。本文把他未强调的半句放大成命题：机器占领练习区的产出后，兑现区相对价格上升；且练习无法完全预装兑现——排练一百次的谈判，开价那一秒仍是首过。

与掌握学习 (Bloom, 1968) 及重测效应文献 (Hausknecht et al., 2007)：前者是重试健身房的宪章，后者是 $\Delta > 0$ 极的测量学证词。本文不推翻它们，只画出适用域边界： $\Delta \geq 0$ 半边。它们对另半边的沉默，正是本文立足之地。

与”认知债务/AI 依赖致退化”论：退化论的轴是”用了会退步”——余额问题，原理可逆，停用再练即可。本文第七节的轴是首过不可补——单遍问题，让渡一次就是全部。混同二者会把结构论证降级为健康忠告；区分二者才能解释为什么”适度使用”这剂退化论药方治不了首过流失。

与赫拉克利特的消解：最后堵一条哲学近路。“人不能两次踏进同一条河”若为真则万物皆单遍，概念因普遍而死。回答：单遍性不是形而上学断言，是测量出来的梯度—— Δ 在重做样本上有分布，多数任务的 Δ 操作上非负（河水换了，但第二次渡河在一切实用意义上是同一件事），少数段落 Δ 显著为负且门位可指认。概念活在对比如里，而对比可计量。

十、讨论：局限、研究议程与证伪条款

10.1 局限

本文的证据链有五处明写的软肋。**代理错位**：错误后果量表偏向后果轴，只能间接代理对象自耗轴——两轴正交论证（4.5节）恰恰意味着这个代理天然含噪。**就业基数**：第二场使用暴露未除以就业规模，职业规模与错误后果若相关则引入偏差。**样本选择**：使用数据来自单一模型的用户，外推受限。**词典粗糙**：第三场的动词词典把语句特征当工序段属性，且判据未防零膨胀——该场判负如实记在这里。**因果识别**：三场皆横截面关联，接管的因果排序有待面板与准实验。

10.2 研究议程： Δ 审计

本文最直接的后续工作是把读数落成数据资产——**跨领域重走落差审计**：从既有登记系统抽取重做样本（插管质量登记、复飞后二次进近的飞行数据监控、资助机构的 A0/A1 档案、重考成绩库、重做手术登记、复职后的留任记录），在可比难度层内估计各域 Δ 及其分布，建立第一张跨域 Δ 表；并建议 O*NET 类任务库增设“重走结构”字段（任务级 Δ 符号与门位 k ），使命题4获得直接检验面——这正是 F3 等待的数据。教育侧则可做推论四的受控检验：配给真实单遍环境的培养组与全重试制对照组，在 $k \leq 2$ 任务上的首过成功率比较。

10.3 证伪条款（写死日期）

F1（承重命题）：若 2031年9月前出现可复核的生产环境证据——AI 系统在无人类首过的 $\Delta \setminus < 0$ 工序（现场谈判、初诊首次接触、危机首响应，任一即可）上成功率不低于该领域人类中位——则“单遍段抗接管”伪，本文作废。**F2（装门命题）**：若到2029年，主要高等教育体系现场口试/监考手写考试占比回落到2022年水平以下——则命题5伪，第六节作废。**F3（接管排序）**：若任务级 Δ 数据建成后显示 AI 实际替代率与任务 Δ 符号无关——则命题4伪。**F4（养成错配）**：若受控比较显示全重试制毕业生与配给单遍环境毕业生在 $k \leq 2$ 任务上首过成功率无差——则第七节错配论伪。**F5（使用侧证据）**：若以就业基数规范化后重跑第二场，负关联消失（偏相关 ≥ 0 ）——则撤回使用侧支持表述，市场判断回落为整体未获裁决。

10.4 清点手册（三处同一口径）

Δ 一律在“同类事件的重做记录”上计量，并在可比难度层内估计或明示选择偏差； k 一律取“该工序的程序文本或操作规程中第一个被标注为不可撤销的步骤的位次”。定义（第四节）、检验（第八节）、证伪（本节）三处同一操作化，不换口径。

十一、结语：给毕业生本人的三句话

理论说完，把答案还给提问里的那个人。

第一句：**盘点你的单遍履历，而不是你的技能清单**。技能清单上每一项都在通胀——AI 会写你会写的、答你会答的。翻出来的应该是另一份账：你首过过什么？哪些事你是在没有重置键的地方做成的——一场真的谈判、一次真的告知、一个真的现场、一回真的收场？给每一条找到它的证人。这份账机器伪造不了，因为它由定义就不接受重投。

第二句：**抢首过，别让渡**。每个值得做的问题只对你敞开一次未被剧透的第一遍。把可重段大方交给机器——检索、初稿、排版、支线上的一切，它做得更好，客气是双输；但在

你打算靠它吃饭的主线上，把首过留给自己，哪怕更慢、更疼、成功率更低。首过的疼是唯一不能外包的学费，也是你账上唯一还在涨价的资产。

第三句：**到装门的地方去报名。**口试、现场、答辩、床边、真客户、真预算——凡是社会重新拧紧门位的地方，就是表现无法伪造、因此人被重新定价的地方。重训健身房里练出的从容在那里帮不了你太多；在那里值钱的是另一种东西：知道这一遍就是唯一一遍，仍然把它走完。

竞争力从来不住在证书里。AI 时代只是把这件事变得无处躲藏：**当一切可以重来的都归了机器，人剩下的、也是人一直真正拥有的，是把只有一遍的事走好的那条命。**

参考文献（分层）

经典层：Münsterberg, H. (1908). **On the Witness Stand**. Doubleday. · Holmes, O. W. (1897). The path of the law. **Harvard Law Review**, 10, 457-478. · Benjamin, W. (1935/1968). The work of art in the age of mechanical reproduction. In **Illuminations**. · Becker, G. S. (1964). **Human Capital**. Columbia UP. · Polanyi, M. (1966). **The Tacit Dimension**. Doubleday. · Bloom, B. S. (1968). Learning for mastery. **Evaluation Comment**, 1(2). · Shackle, G. L. S. (1949). **Expectation in Economics**; (1961). **Decision, Order and Time in Human Affairs**. Cambridge UP. · Goffman, E. (1967). Where the action is. In **Interaction Ritual**. Anchor. · Arendt, H. (1958). **The Human Condition**. University of Chicago Press. · Spence, M. (1973). Job market signaling. **Quarterly Journal of Economics**, 87(3), 355-374. · Arrow, K. J., & Fisher, A. C. (1974). Environmental preservation, uncertainty, and irreversibility. **QJE**, 88(2), 312-319. · Pocock, S. J. (1977). Group sequential methods in the design and analysis of clinical trials. **Biometrika**, 64(2), 191-199. · O'Brien, P. C., & Fleming, T. R. (1979). A multiple testing procedure for clinical trials. **Biometrics**, 35(3), 549-556.

专著层：Dixit, A. K., & Pindyck, R. S. (1994). **Investment under Uncertainty**. Princeton UP. · Phelan, P. (1993). **Unmarked: The Politics of Performance**. Routledge. · Auslander, P. (1999). **Liveness: Performance in a Mediatized Culture**. Routledge. · Gann, G. D., et al. (2019). International principles and standards for the practice of ecological restoration (2nd ed.). **Restoration Ecology**, 27(S1), S1-S46.

论文层：Prahalad, C. K., & Hamel, G. (1990). The core competence of the corporation. **Harvard Business Review**, 68(3), 79-91. · Autor, D. H., Levy, F., & Murnane, R. J. (2003). The skill content of recent technological change. **QJE**, 118(4), 1279-1333. · Autor, D. H. (2015). Why are there still so many jobs? **Journal of Economic Perspectives**, 29(3), 3-30. · Frey, C. B., & Osborne, M. A. (2017). The future of employment. **Technological Forecasting and Social Change**, 114, 254-280. · Brynjolfsson, E., & Mitchell, T. (2017). What can machine learning do? Workforce implications. **Science**, 358(6370), 1530-1534. · Acemoglu, D., & Restrepo, P. (2019). Automation and new tasks. **Journal of*

Economic Perspectives*, 33(2), 3–30. · Bradshaw, A. D. (1987). Restoration: An acid test for ecology. In Jordan, Gilpin & Aber (Eds.), **Restoration Ecology**. Cambridge UP. · Benayas, J. M. R., Newton, A. C., Diaz, A., & Bullock, J. M. (2009). Enhancement of biodiversity and ecosystem services by ecological restoration: A meta-analysis. **Science**, 325, 1121–1124. · Moreno-Mateos, D., Power, M. E., Comín, F. A., & Yockteng, R. (2012). Structural and functional loss in restored wetland ecosystems. **PLOS Biology**, 10(1), e1001247. · Moreno-Mateos, D., et al. (2017). Anthropogenic ecosystem disturbance and the recovery debt. **Nature Communications**, 8, 14163. · Suding, K. N., Gross, K. L., & Houseman, G. R. (2004). Alternative states and positive feedbacks in restoration ecology. **TREE**, 19(1), 46–53. · Hobbs, R. J., Higgs, E., & Harris, J. A. (2009). Novel ecosystems. **TREE**, 24(11), 599–605. · Murcia, C., et al. (2014). A critique of the “novel ecosystem” concept. **TREE**, 29(10), 548–553. · Fukami, T. (2015). Historical contingency in community assembly. **AREES**, 46, 1–23. · Loftus, E. F., & Palmer, J. C. (1974). Reconstruction of automobile destruction. **JVLVB**, 13, 585–589. · Loftus, E. F., Miller, D. G., & Burns, H. J. (1978). Semantic integration of verbal information into a visual memory. **JEP: HLM**, 4(1), 19–31. · Wells, G. L. (1978). Applied eyewitness-testimony research. **JPSP**, 36(12), 1546–1557. · Loftus, E. F. (1993). The reality of repressed memories. **American Psychologist**, 48(5), 518–537. · Nader, K., Schafe, G. E., & LeDoux, J. E. (2000). Fear memories require protein synthesis in the amygdala for reconsolidation after retrieval. **Nature**, 406, 722–726. · Deffenbacher, K. A., Bornstein, B. H., & Penrod, S. D. (2006). Mugshot exposure effects. **Law and Human Behavior**, 30(3), 287–307. · Calabresi, G., & Melamed, A. D. (1972). Property rules, liability rules, and inalienability: One view of the cathedral. **Harvard Law Review**, 85(6), 1089–1128. · Kronman, A. T. (1978). Specific performance. **University of Chicago Law Review**, 45(2), 351–382. · Schwartz, A. (1979). The case for specific performance. **Yale Law Journal**, 89(2), 271–306. · Radin, M. J. (1987). Market-inalienability. **Harvard Law Review**, 100(8), 1849–1937. · Sakles, J. C., et al. (2013). The importance of first pass success when performing orotracheal intubation in the emergency department. **Academic Emergency Medicine**, 20(1), 71–78. · Hausknecht, J. P., et al. (2007). Retesting in selection: A meta-analysis. **Journal of Applied Psychology**, 92(2), 373–385. · Ericsson, K. A., Krampe, R. T., & Tesch-Römer, C. (1993). The role of deliberate practice in the acquisition of expert performance. **Psychological Review**, 100(3), 363–406. · Peters, O. (2019). The ergodicity problem in economics. **Nature Physics**, 15, 1216–1221. · Gelman, A., & Loken, E. (2014). The statistical crisis in science. **American Scientist**, 102(6), 460–465. · Felten, E., Raj, M., & Seamans, R. (2021). Occupational, industry, and geographic exposure to artificial intelligence. **Strategic Management Journal**, 42(12), 2195–2217.

近文层：*Metis AI* (arXiv:2605.14407, 2026) ——“后果不可逆性”支柱的当代占位表述。 · Eloundou, T., Manning, S., Mishkin, P., & Rock, D. (2023). GPTs are

GPTs. arXiv:2303.10130. · Anthropic Economic Index (2025–2026 各期) ——基于真实对话的职业/任务级使用测量。 · 2023–2026 年各国大学恢复口试与现场手写考试的政策报道（装门回潮的实证层）。

数据与预注册：O*NET Database 30.0 (Work Context 4.C.3.a.1 ; Job Zones ; Work Activities 4.A.3.a.1) , U.S. Department of Labor. · AIOE 数据 (github.com/AIOE-Data/AIOE) . · Anthropic Economic Index 职业级使用暴露与任务级渗透表 (huggingface.co/datasets/Anthropic/EconomicIndex) . · 预注册一 SHA-256 : 5410b75bb46768cb928f188d60c992e0df33dde0aa0bb96fe7d72fdd287916b; 预注册二: 780a733c23857125359ae6b74d945e90138eece91353dd6887022685911eec35; 预注册三: 7c63f87312605e185085bd25c583ac10515a01d08e524fe9bdc09ce3220ab1fa (均先于对应数据下载, 2026-09-02) 。

第七章 现配优势

【章首】本章给方程添的第三条本体性质：**当场性**——竞争差的载体不是任何可交接的存量，而是每次交手由各自平庸的可传构件加一手不入账的配组当场做出的显露。四条否定性质一口气：照单交接会回落，逐项拆开不显值，公开出去也偷不走，写进清单那一层就立刻被批量供给。读数是留种率 s ：把优势写成尽可能完备的交接清单交给同资历者照做，接收方能复现的份额——人类版的杂交二代。它与《常席差》共锚“配组”而各管一头：本章管优势的载体（交接后复现多少），常席差管读数的归属（残差记给谁）；同一人清单复现率高而跨配组残差亦高，两量即证独立。与《入账之前的技能》的咬合也值得先记：本章“写进清单即被批量供给”与彼章“文件载体入账即入可复制区”是同一现象的两读。本章欠的账：同质多案例的交接复现实验（ $N \geq 6$ ），未做。

大白话：本事不在菜谱里，在配菜的那一手上

一个厨师做出一桌好席。你把他的菜谱全要来了——每一道的用料、火候、步骤，一字不差。你照着做，做出来的东西能吃，但就是不对。

问题出在哪？不在菜谱漏了什么。在于那一桌席不是菜谱的总和，而是他当场根据这批食材、这个灶、这些客人、这个季节配出来的。配的这一手，写不进菜谱——不是他藏私，是它本来就不是可以写下来的东西：它每次都不一样，因为每次的料和人都不一样。

《现配优势》讲的就是这件事：**竞争差的载体，不是任何可以交接的存量，而是每次交手当场配组做出来的显露。**

它有四条听起来很怪、但都成立的性质：

一，照单交接会回落。你把自己的做法写成尽可能完备的清单交给同资历的人，他照做，效果会掉一截。掉的那一截，就是当场配组的那一份。育种上有个现成的类比：杂交种长得极好，但你留它的种子再种，第二代就退化了——好处不在种子里，在那次杂交里。

二，逐项拆开不显值。把这一手拆成一条条能力，每一条看起来都很平庸：会沟通、懂数据、了解客户。平庸的构件配出不平庸的东西，这正是“配”的意思。

三，公开出去也偷不走。因为拿走的是清单，不是配组。

四，写进清单那一层，立刻被批量供给。这条最狠：任何一项能力，一旦被清楚地命名、写进课程和模板，它就会被大批量地供应出来，于是它的差异度迅速归零。**能力清单是一个自毁装置**——写得越清楚，失效得越快。

为什么今天要紧？

因为 AI 是有史以来最强的“清单执行器”。凡是能被说清楚的做法，它都能立刻大规模照做。于是所有可陈述的能力都在贬值，剩下的价值集中到了那一手不可陈述的配组上。这不是玄学：它可以被测。

读数是什么？

叫**留种率 s** ：把你的优势写成一份尽可能完备的交接清单，交给一位同资历的人去执行，看他能复现多少。复现得越多，说明你的优势越是存量（也就越快会被复制）；复现得越少，说明它越是当场配出来的。

请注意这是一个反直觉的指标：**留种率低是好消息**。一个交接清单能被完美复现的人，正在被自己的清单替换掉。

一个小场景。

一家公司的金牌销售被要求"把方法论沉淀下来"。他认真写了三十页 SOP。半年后，新人照着做，业绩只有他的四成。管理层的结论通常是"SOP 写得不够细"，于是要求写六十页。

真实情况可能是：他的本事就不在 SOP 里——他每次进客户办公室的头三分钟，是根据当天的气氛、对方桌上的东西、上一次会议留下的尴尬，**当场配出来的**。写到一百页也没用，那一手不在纸上。正确的动作不是继续写细，是**让新人和他一起出现在现场，配组给他看**。

再看两个身边的例子。

带孩子。育儿书写得再细，也没法告诉你此刻这个孩子、在这个时间、因为这件事哭，你该怎么办。有经验的家长做的事，是当场把"孩子的状态、自己的耐心余额、今天的日程、这件事的严重程度"配在一起，做出一个此刻的处理。育儿书是构件，那一手是配组。这也解释了为什么"照书养"经常不灵——不是书错了，是书里没有、也不可能有一手。

主持一场饭局。谁和谁不能坐一起、话题什么时候该转、谁需要被点名说两句、什么时候该结束——这些没有一条能写成流程，但做得好和做得差之间的差别，所有在场的人都感觉得到。这是纯粹的当场配组，而且它几乎完全不可传授，只能通过多主持几次长出来。

三种常见的误读。

第一种：以为它是"经验"。经验是存量，是可以积累、可以陈述、可以传授的东西。当场配组不是存量：它每次调用的是不同的构件组合，而组合方式取决于当下的条件。**有一个有二十年经验但二十年只在同一种条件下工作的人，配组能力可能不如一个换过五种条件的年轻人**。

第二种：以为写不清楚就是没方法。恰恰相反：一个人如果能把自己的优势完全写清楚，那说明他的优势已经变成了存量，而存量会被批量供给。**写不清楚不是缺陷，是那份价值还活着的证据**。

第三种：把"不可陈述"当成不可训练。这是最有害的一种，因为它让人放弃练习。不可陈述只意味着不能靠"讲"来教，但完全可以靠**安排练习的机会**来长——换搭档、换约束、换场景，一次比一次不同。所有配组能力都是这么长出来的，没有例外。

给自己做个小检查。如果明天你要把手上的活完整交给一位同资历的人，你需要写多少页？写完之后，你估计他能做到你的几成？那个"几成"就是你的留种率——它越低，你越安全。

它和旁边几个概念的区别。

和《常席差》共用一个词"配组"，但管的是两头：这一章管**载体**——你的优势交接出去还剩几成；《常席差》管**归属**——这次成果里哪一份该记在你名下。一个人完全可能交接清单复现率很低（说明本事在配组里），同时跨配组残差也很高（说明那份配组能力稳定跟着他）。两个量各测一头。

和《零判份》的区别在于比较对象：这一章比的是"照单执行的人能复现多少"，《零判份》比的是"世界现存样本能不能拿出同类物"。前者是**人对人**，后者是**你对世界**。同一份产

出，留种率可以很低而未见率也很低——你的做法别人学不会，但世界早有别的路径做出同样的东西。

和《入账之前的技能》咬合得很紧：这一章说“写进清单那一层就被批量供给”，那一章说“写进文件载体就进了可复制区”——**同一件事的两种读法**，一个从竞争差看，一个从归属看。

一句话记住它：这一章管的是**装在哪里**——你的本事装在可交接的东西里，还是装在场那一手里。

怎么长出来：方法、技术与流程

个人层面。

别把力气全用在“把能力写清楚”上，那是自我批量化。真正该做的是**增加配组次数与配组多样性**：换合作对象、换场景、换约束条件，逼自己每次重新配一遍。一个人的当场配组能力，几乎完全是被“次数×多样性”喂出来的。

大学发生学教育：把“配”变成可练的东西。

现行教育的默认动作与此完全相反——它把一切拆成可教的模块，逐个考核。模块化是必要的，但如果四年里学生只做模块，他毕业时手上全是构件，没有一次配的经验。

- **轮换配组制。**项目类课程里，**同一个学生一学期必须进入至少三个不同构成的小组**，且组员由抽签决定，不许自由组队。自由组队会让人固定在同一批熟人里，配组次数看着很多，实际只有一种。
- **交接复现作业（这是本章最有力的一个教学设计）。**学生 A 完成一个任务后，写一份“完备交接清单”交给学生 B；B 照单执行，第三方评分。然后当堂公布两人分差——那个分差就是 A 的当场配组的可见证据。这个作业一次就能让学生明白“我的本事在哪里”，比讲十次都管用。
- **分离评分。**同一个作业给两个分数：**清单件分**（按清单执行得到的部分）与**配组件分**（超出清单的部分）。两个分数分开记，学生第一次看到它们不是一回事。
- **约束抽签。**同一个题目，不同小组抽到不同的硬约束（预算减半、少一个人、时间压缩、必须用某种材料）。这一步是配组训练的核心：**配的本事只在约束变化时才被调用。**

教师的动作。

把“标准答案示范”改为“当场配组示范”：教师拿一个自己也没准备过的现场条件，当着学生的面配一次，包括配砸。学生需要看见的不是完美流程，是在**不完美条件下重新配的过程**。

最容易做坏的地方。

把“配组能力”写成一条新的能力清单去教——那会立刻触发第四条性质，它写清楚的当天就开始被批量供给。配组只能被练，不能被讲；本节所有设计的重点都在安排练的机会，而不在定义它是什么。

原文

现配优势

——**人工智能时代大学毕业生竞争差的载体结构与留种率**

摘要 本文提出并论证一个关于个人竞争优势之载体结构的判断：在可陈述认知劳动的复制成本迅速趋近于零的条件下，一名大学毕业生身上真正构成竞争差的部分，既不是一份可公示、可随身携带的结果存量，也不是一套可以照单传承的能力配方，更不是一件靠不公开而圈住的独知资产，而是在每一次具体交手中，由若干各自平庸的可传构件加上一手不入账的配组，当场重新做出来的显露。本文称之为**现配优势**（freshly-assembled advantage），并给出与之配套的可清点读数——**留种率 s**：把一项优势写成尽可能完备的交接清单，交给同资历者照做，接收方所能复现的份额。现配优势的四条否定性质是：照单交接会回落，逐项拆开不显值，公开出去也偷不走，写进清单那一层就立刻被批量供给。本文用一个二乘二的辨别格把它与通用技能、制度化流程、默会手艺三者分开，指出当前就业力话语的主流答案恰好落在一个特定的失效格中——本文称之为**满单件**：清单全勾、认证齐全、短期指标改善，照做却回到均值，而且失效不产生任何信号。本文对若干既有理论（贝克尔的专用人力资本、波兰尼的默会知识、斯彭斯的信号理论、资源基础观、格罗伊斯伯格的明星可携带性研究、工人—企业固定效应分解、个人商誉估值实务）逐条给出判定形式，说明在何种观察下本文为对而它们为错，反之亦然。本文报告一次事先登记的经验检验，其口径、判负条款与停止规则均在读取任何数据之前写死：以美国全职开发者的公开薪酬调查为对象，控制从业年限与学历后，“可列举技能条目数”对报酬的边际关联由 2019 年的每项 +0.24%（不显著）变为 2024 年的 -0.50%（显著）。本文同时报告两处对本文原判断构成实质削弱的反向材料，并据此收窄了主张的适用范围。最后给出证伪条款、清点手册、反噬预言与伦理限制。

关键词 竞争优势的载体；留种率；可清单化；配组；就业力；人力资本

一、两份一模一样的简历

设想一位招聘者手上有两份简历。学校同档，绩点相差不到零点一，实习经历数量相同，掌握的工具与语言逐条对齐，证书目录几乎可以互相替换，作品集的完成度都过得去。招聘者知道这两个人不一样，也知道其中一个日后会明显强过另一个——他只是在这两张纸上找不到那处不一样。

这是一个古老的困境，但它在最近几年换了性质。过去的困境是**信息不足**：那处不一样确实存在于某个人身上，只是没被记录下来，所以要靠面试、试用、推荐信去把它逼出来。今天的困境是**信息不足的那一面正在被填满，而困境反而更深了**。今天一名毕业生能够陈述、演示、附上作品链接的一切，几乎都可以被一个足够熟练的使用者在半小时内用公开工具做出功能等价物。可陈述的那一层不是没被记录，而是记录得太好、太便宜、太容易被复制，以至于记录本身不再起区分作用。

于是问题就从“如何把那处不一样记录下来”变成了一个更难的问题：**如果一切能被记录的都不再区分人，那么区分人的东西究竟以什么方式存在？**

本文的回答不诉诸“更高阶的能力”。这是一条已经被走烂的路：每当某一层技能贬值，主流回答就往上提一层——从具体工具提到编程能力，从编程能力提到分析性思维，从分析性思维提到批判性思维、沟通力、韧性、创造力。这条路的问题不在于所提之物不重要，而在于它的**贬值机制与被它取代的东西完全相同**：一项能力一旦被命名并写上榜单，课程、教材、培训与提示词模板会在一到两年内跟上，把它变成一件人人可陈述的物件，于是它重新掉回原处。世界经济论坛《未来就业报告》二〇二五年版把分析性思维列为最核心的技能，约七成雇主视之为必备。这份榜单如实反映了雇主的需求，但它同时也是一台把稀缺变成普及的机器：任何上榜的能力，其上榜本身就是它开始失去区分力的信号。

本文认为，问题出在提问的方式上。所有这些回答——不论所提之物有多高阶——都共享同一个未被检查的假定：**竞争优势是一种可以有明确归属的东西，它先于每一次具体交手就已经存在于某个持有者身上，交手不过是把它读出来。**本文要论证的是，正是这个假定在今天失效了；而一旦把它拆掉，"AI 时代毕业生的核心竞争力是什么"这个问题就有了一个可清点、可证伪、且与经验相符的回答。

在进入论证之前，先给三个日常景象，它们各自已经包含了本文的一个要点。

其一，杂交稻不能自留种。买来的杂交种子长势齐整、产量高，但把它结的籽留到第二年种下去，产量会掉下来，植株参差不齐。玉米单交种的自交衰退实测大约在百分之十到百分之三十几之间。更值得注意的是第二件事：配出这个高产杂交种的那两个亲本自交系，各自单独种下去，长势比普通品种还差。优势不长在任何一粒种子身上，也不长在任何一个亲本身上，它长在"哪两个配哪两个"这一手里。

其二，盘一家老面馆要付转让费。桌椅、锅灶、冰柜、招牌可以按件点清，可是接手的人还得多付一笔说不清明细的钱。那笔钱买的是什么？不是任何一件可以搬走的东西，而是这家店在这条街上、和这批老顾客之间、由这套做法维持着的那个状态。它列不进任何一张清单，却在整体成交的那一刻被明确定价。

其三，一对配合多年的双打搭档拆对以后，两个人的成绩都会掉。这件事在体育里几乎是常识，从来没有人认为拆对的那一刻"有一项能力从某个人身上消失了"。但如果我们相信能力是长在个人身上的存量，那就必须解释：为什么把两个存量分开摆放，总量会变小？

这三个景象说的是同一件事：**有一类优势，它的载体不是某个持有者，而是一个每次都要重新做出来的配组。**本文要做的，是把这句常识变成一个能进考核、能被证伪、能对教育与雇用提出具体建议的判断。

二、三种既有回答，以及它们共有的那块地基

关于"一个人有多强"这件事，学术与实务界有三条各自成熟、彼此对立的回答路线。它们的对立是真实的；但它们脚下是同一块地。

第一条路线：看结果台账。这条路线的纯粹形态出现在医学里。一九一四年前后，波士顿外科医生欧内斯特·科德曼提出"最终结果"制度：每一家医院都应当追踪每一个病人的最终结局，并把它公开，让人们据此判断谁做得好、谁做得差。这个主张后来长成了两大分支：一是绩效公示运动——风险校正死亡率、并发症率、再入院率被逐年公布；二是量一效应研究——做得多的医院与医生结果更好，于是"手术量"成为能力的代理指标。这条路线对能力的看法极其干脆：能力就是台账上那个经过风险校正的数字，别的都是自我陈述。它在教育领域的对应物是绩点、标化成绩、竞赛名次、作品集与"技能矩阵"。

这条路线内部也有一个强有力的批评派。批评者指出，公示会诱导选择效应：外科医生开始回避高风险病人，数字变好而病人变差。这场争论持续了三十年，但值得注意的是，争论双方都同意一件事：**台账所记的那个数，是那位医生本人的一项属性**；分歧只在于它测得准不准、会不会被操纵。

第二条路线：看配组路径。这条路线的纯粹形态出现在育种学里。一九〇八年，乔治·沙尔在玉米上系统地记录了杂种优势：把两个各自衰弱的自交系杂交，后代显著超过双亲。一九四二年，斯普拉格与塔特姆给出了沿用至今的分解：**一般配合力**是一个自交系在各种杂交组合中的平均表现，**特殊配合力**则是某些特定组合的表现好于或差于按各自一般配合力所预期的那些情形。这个分解的意思是残酷的：**一个亲本自己的长相，预测不了它配出来的结**

果；有一部分优势只在特定的配对里出现，不在任何一方身上。而且这份优势不能自留——把杂交种的后代留作种子，优势就散掉了。

这条路线对能力的看法是：能力是配组的属性，不是构件的属性；要知道两个构件配起来行不行，只有真配一次。

第三条路线：看条件场。这条路线的纯粹形态出现在商业秘密法里。一件东西要被承认为商业秘密，须同时满足三个要件：不为公众所知悉、持有人采取了合理的保密措施、并且**正因为不为人所知而具有价值**。一九七四年，美国最高法院在 Kewanee Oil 案中确认各州商业秘密法与联邦专利法并行不悖，而该案中的禁令措辞极能说明问题：保护持续到该秘密被公开之时为止。独立研发合法，反向工程合法，一旦公开则权利消灭且无从恢复。

这条路线对能力的看法是：价值不在物本身，而在物所处的那个“别人还没有”的条件场里；条件一变，同一件物就一钱不值。它在就业市场里的对应物是稀缺性叙事：你的价值取决于市场上还有多少人能做同样的事。

这三条路线彼此为敌，而且敌意是真的。公示运动每天的日课——把结果台账摊开给所有人看——按商业秘密法的定义恰恰是资产的自毁行为；反过来，商业秘密所要维护的“不公开”，正是公示运动所诊断的病灶。育种学则对另外两家同时翻脸：它告诉台账派，照着最好的显露去复制必然失败，因为显露里不含配方；它又告诉商密派，我把产物本身公开出售——种子卖的就是杂交种子本身——优势照样不泄露，因为配方不在产品里。

但请注意它们脚下那块共同的地：**这三家争的是这份优势应当记在哪一本账上**——记在个人的结果台账上，记在配组的谱系上，还是记在条件场的资格上。它们从未争过一件事：**这份优势是一份先于每次交手就已经存在、可以归到某一个单一持有者名下的存量，而每一次交手只是把它读出来。**把这句话念给三家中的任何一家听，得到的回答都会是“这还用说吗”。

本文把这个未被争论的假定命名为**单一载体假定**：优势必须由某一个单一的载体持有——个人、团队、信号、资产、组织——差别只在于是哪一个。下一节要说明，这块地基上有一道裂缝，而且这道裂缝是由第一条路线自己挖出来的。

三、来自台账派自己的一次观察

二〇〇六年，赫克曼与皮萨诺在《管理科学》上发表了一项研究。他们利用同一个州内多家医院的手术记录，追踪了那些同期在不止一家医院执业的心脏外科医生。这是一个非常干净的设计：同一个人、同一段时间、同样的资历与训练，唯一变化的是他此刻站在哪家医院的手术台前。

结果是：这些医生的手术结果质量，随着他**在这一家医院的近期手术量**上升而改善，却不随他**在其他医院的手术量**改善。也就是说，一个外科医生在 A 院积累的表现优势，并不能整体地跟着他走进 B 院的手术室。

这项研究在自己的领域里回答的是一个组织学习的问题：绩效在多大程度上是可携带的？它被归入“人力资本的组织专用性”这一议题，通常被读作对贝克尔式“企业专用人力资本”的一次实证支持——某些技能只在特定组织里才值钱。这是一个合理的读法，但它把一件更锋利的东西钝化了。

请回到台账派自己的信念上来看这项观察。台账派主张：那个经过风险校正的数字，是这位医生本人的属性。如果真是如此，这项属性没有理由在他换一张手术台时失灵——它是他的，不是那张台子的。而观察到的事实是：**读数与“他在这里最近做了多少”绑定，与“他**

这一生总共做了多少"不绑定。一份长期积累的、写在履历上的总量，并不能保证今天这一台的表现；而一段近期的、在这个具体场所内的密集配合，能。

这意味着那个数不是从这位医生身上"读出来"的，而是由"这位医生 × 这家医院"这一对配组在最近这段时间里制作出来的。台账没有在记录一份存量，它在记录一次次制作的产物，然后把产物归到了其中一个构件名下。

这就是本文的支点，而它有三个性质使它足够可靠。**第一，它来自这三条路线中最主张存量的那一家**，是台账派自己的实证观察，不是外人的批评。**第二，它是一个观察结果而非一个主张**——它不是有人论证说能力不可携带，而是数据显示它没有跟着人走。**第三，它没有被用在这个方向上**：三十年来它被读作"某些技能是组织专用的"，也就是仍然是一份存量，只不过归属的名下从个人改成了个人与组织的匹配；它从未被读作"存量先在"这条前提本身的反证。

另外两家其实各自也知道一件同类的事实，只是同样不往这个方向用。育种学知道：**特殊配合力只有真配一次才测得出来**，亲本各自的表现预测不了组合的表现——这句话如果照字面理解，说的就是那份优势在配组发生之前并不存在于任何地方。商业秘密法知道：**它所保护的从来不是那件物本身，而是"尚未被公开"这个状态**——独立研发合法、反向工程合法、公开即权利消灭，这三条加在一起说的是，被保护的对象是一个条件在场的状态，而不是一份可以放进保险柜的存量。

三家各自握着一块碎片，而三块碎片拼起来正好把三家共同站着的那块地抽掉。

四、承重命题

本文的承重命题是：

在可陈述认知劳动的复制成本趋近于零的条件下，一名大学毕业生身上构成竞争差的部分，**不是**台账上可公示、随人携带的结果存量，**不是**可照单传承的能力配方，**也不是**靠不公开而圈住的独知资产，**而是每一次交手中由若干各自平庸的可传构件加一手不入账的配组当场重新做出来的显露。**

这个命题的三重否定分别对准前述三条路线各自的划界处，而它的正面内容有四条可检查的性质：

性质一：不可自留。把这份优势写成一份尽可能完备的交接清单，交给一位同资历者照做，接收方的表现会明显回落到该资历段的均值附近。这不是因为清单写得不好，而是因为清单里能写下的都是构件，配组那一手写不进去。

性质二：不可逐项计价。把这份优势拆成构件逐项估值，每一项都平庸，加总起来小于整体。它只有在整体成交的那一刻，才以"说不清明细的那一笔"被定价——盘店的转让费、公司并购中的商誉、球队买断整条防线时那笔溢价，都是同一个位置上的同一笔钱。

性质三：不因公开而失效。你可以把作品全部公开、把方法写成教程、把过程录成视频，别人拿走的仍然只是"种子"，配不出你的下一茬。这一条与商业秘密的逻辑正相反：现配优势不靠不公开来维持，因此也不会因公开而消灭。

性质四：不随人整体携带。换一个搭档、换一个单位、换一批客户，这份优势不会整体跟过去，需要新的场里重新做一遍。它不是行李，是一次次的现做。

本文把这个存在物命名为**现配优势**：现配，指它在每次交手中当场配成；优势，指它确实构成人与人之间可观察的差别。它不是"隐性能力"（隐性能力仍然是长在人身上的存量，

只是没被记录)，也不是"运气"（运气不可复现，而现配优势在同一配组内高度可复现），也不是"人脉"（人脉是关系的存量，现配优势是关系被使用时的产物）。

下一节给出使这个命题可以被清点、从而可以被检验的那个读数。

五、读数：留种率 s

一个关于载体结构的判断，如果不能落到一个可以被清点的数上，就只是一种说法。本文给出的读数是**留种率 s** 。

定义。取一位持有者在某一类任务上最近一次被承认的表现读数 R 。请他把这项工作写成一份尽可能完备的交接清单——步骤、判断依据、常见陷阱、模板、检查表、联系人、可复用的产物，凡是能写下来的都写。把这份清单交给一位同资历者，让他在同类任务上独立执行。取接收方前三次执行的读数均值 R' 。则

$$\text{留种率 } s = R' / R$$

s 接近一，说明这项优势可以自留：写下来就传得下去。 s 显著小于一，说明这项优势的相当一部分不在清单里，也不在被交出的构件里。

这个读数的形式借自育种学的自交衰退比：杂交第一代与其自交后代的产量之比。这不是修辞。两处的结构是同一个：**把一次配组的产物当作种子再种一次，看还剩多少。**

清点手册。为避免读数在使用中漂移，本文把口径写死一处，正文、检验、赌注三处一律使用这一处口径，不另设变体：

1. **同类任务**指同一客户或同一题域下的同一类交付。换客户、换题域一律记"不适用"，不折算。
2. **同资历者**指在公开可核的资历维度（学历、年限、职级、相关证书）上与原持有者处于同一档的人。资历不同一律记"不适用"。
3. **完备清单**的判定不由作者自评，而由第三方按预先约定的条目框架检查；缺项超过一成记"不适用"。
4. **读数 R** 必须取自组织本来就在记录的指标，不得为测 s 而新造指标；没有既有台账的，记"无从判定"，**不记零**。
5. **前三次执行取均值**，是为了容纳一次性的适应成本；若组织只允许一次执行，记"单次读数"并单独标注。
6. s 的取值不设上界。接收方超过原持有者（ s 大于一）是有信息的结果，说明原读数中含有场所或时点的成分而非配组的成分。
7. **s 是位置的属性，不是人的属性**：它描述的是"这项优势有多大比例能经清单转移"，不是"这个人有多强"。

这七条构成一个可以被别人重复的清点程序。凡本文后文提到 s ，皆指此程序的产物。

为什么用"能不能传下去"来测"是不是存量"。这一步是本文推理的关节，需要说清楚。如果一份优势确实是一份先在的、归属于某个持有者的存量，那么把它形式化、写下来、交出去，损失的应当只是形式化的精度——写得越细，损失越小， s 应当随清单完备度单调趋近于一。反过来，如果一份优势是每次由配组现做出来的，那么无论清单写得多么多，被交出

去的都只是构件，而做出优势的那一手在新的配组里必须重新长出来，s 就会顽固地停在一个明显小于一的地方，并且**不随清单完备度改善**。

于是"s 是否随清单完备度趋近于一"就成了一个可以把两种载体结构分开的判据。这不是定义上的循环，因为清单完备度是可以独立于 s 被测量的。

六、辨别格

有了 s，就可以把"一名毕业生身上的各种优势"分类。用两个维度：**这项优势能否被写成完备的交接清单（可清单化），以及它的留种率是高是低**。

	高留种率（s 接近 1）	低留种率（s 明显小于 1）
可清单化	① 商品件 ：教材技能、证书内容、标准工具的使用。照单复现成功，因此市场与自动化会大量供给它。复现即无差。	② 满单件 （本文靶格）：被课程化的高阶能力。清单全勾、认证齐全，照做却回到均值。
不可清单化	③ 默会手艺 ：师徒身传、随人迁移的手感与判断。写不下来，但换个地方仍然管用。	④ 现配区 ：优势在配组里，既列不全也带不走，只能当次做出。

四个格各有其现实对应物，也各有其命运。

格①商品件。它是可清单化的，也确实能传下去——所以它一定会被普及。这不是坏事，正相反，一个社会的商品件越丰厚，其从业者的地板越高。但正因为它能被复现，它就不能构成竞争差：**凡是照单能做成的，早晚人人做得成**。今天的新变化只是"早晚"变成了"随即"。

格③默会手艺。波兰尼所说的"我们知道的比我们能说出来的多"落在这一格：写不下来，但它随人走。老师傅换一家工厂，手感还在。本文不与这一格争地盘——它是真实存在的，而且在很多行业里仍然是竞争差的主要来源。本文只指出一点：默会手艺的**转移方式是人的整体迁移**（挖人、拜师、长期共处），因此它的经济学与格④完全不同——格③挖人有效，格④挖单个人无效。

格②满单件。这是本文的靶格，下一节专论。

格④现配区。它既写不全，也带不走。它的存在方式是：一批各自平庸的构件（其中许多本身就是格①的商品件），加上一手在特定人、特定场、特定问题之间的配法，在交手当场做出一个别人照单复现不了的结果。杂交种的产量在这一格，双打组合的成绩在这一格，那家老面馆值转让费的部分在这一格。

四格的动态。在可陈述劳动复制成本趋零的条件下，格①正在急剧扩容并贬值；格②在扩容的同时保持着虚假的价值感；格③相对稳定但传播缓慢；格④是唯一一个其价值随格①扩容而**上升**的格——因为构件越便宜、越普及，"能把便宜构件配出不便宜结果的那一手"就越稀缺。

这就是本文对开篇问题的直接回答：**AI 时代大学毕业生的核心竞争力，在格④**。

七、靶格：满单件

一个理论如果只说"什么是对的"，它就没有承重能力。它必须能指出一个具体的、正在被大量生产的、失败而不自知的东西。本文指出的是格②：**满单件**。

满单件是这样一类东西：它被写成了完备的清单，接收方也照单执行成功，形式审查全部通过，短期指标全面改善——但它的留种率高得恰恰证明它不构成竞争差。**它的成功就是它的失效**。

满单件的三个签名：

签名一：短期指标改善。引入某项能力培养方案之后，可观测的指标立刻好转：课程完成率、认证获取数、作品集提交量、面试通过率、雇主满意度问卷的分数。所有这些改善都是真的。

签名二：失效不产生信号。关键在这里。当一项能力被普及之后，它不再区分人；但"不再区分人"这件事不会在任何一张报表上显示为一个坏数。持有者没有做错任何事，清单每一项都勾了，交付质量也没有下降。失效以"大家都一样好"的形式出现，而"大家都一样好"在现有的评价体系里是成功的样子。这就是为什么满单件能持续多年不被发现。

签名三：自我加固。由于失效不产生信号，系统对"没有拉开差距"的唯一反应是**加长清单**：再加一项能力、再加一门课、再加一个认证、再加一层评估。而每一次加长都在生产新的满单件。清单越长，格②越肥，格④越被挤压——因为准备清单要花时间，而那些时间本可以用来真刀真枪地和具体的人配一次。

榜单机制正是满单件的生产线。一项能力被识别为稀缺、被写上榜单、被课程化、被认证化、被写进职位描述，这条路径的每一步都是善意且理性的，而它的终点是这项能力不再区分人。**上榜是失效的开始，而不是价值的证明。**世界经济论坛的技能榜每两年更新一次，榜单上的名词不断向上迁移，这个迁移速度本身就是格②扩容速度的一个粗略读数。

需要说清楚本文**不主张**什么。本文**不主张**这些能力不重要。分析性思维当然重要，沟通当然重要。本文主张的是关于**竞争差**的一个更窄的命题：一项能力的重要性与它的区分力是两回事，而教育与雇用体系长期把两者混为一谈。呼吸对生存极其重要，但没有人靠呼吸得分。

八、五个场景

一个判据如果只在纸上成立，它就是修辞。本文给出的判据是一句不含任何情态词、可以在真实材料里查的问题：

把这项优势写成完备的交接清单交给同资历者照做，接收方的读数落在原表现的几成？这个数记在哪份正式记录里？

下面五个场景，答案各不相同；其中四个的答案是可以在公开材料里查到的。

场景一：杂交种自留。答案：明显低于一。玉米单交种的自交后代产量下降幅度在多个实测系列中大约落在百分之十到百分之三十几之间。记录在哪里：在产量对比试验的正式报告里，也在每一袋种子的包装说明上——"不可留种"是写在商品上的规则。**这是一个 s 明显小于一而且被制度化承认的案例。**

场景二：世界卫生组织手术安全核查表。答案：接近一，甚至大于一。这份不到二十项的清单被交给条件差异极大的医院照做，术后并发症与死亡率显著下降，效果可复现，随后在全球范围内被采纳。记录在哪里：多国多中心的前后对照研究与随后的推广政策文件里。**这是一个 s 接近一的成功案例，而且正因为成功，它在十年之内变成了地板——今天没有任何一家医院能靠"我们用核查表"构成竞争优势。**这个场景对本文极其重要，第十节将专门处理它对本文的削弱作用。

场景三：名店菜谱出版。答案：介于两者之间，且长期停在那里。菜谱写得越来越细——克数、火候、时间、图解、视频——家庭复现者做出来的东西"对，但不是那个味"，而

本店照旧排队。记录在哪里：没有正式记录。这正是本文要指出的问题之一：**格④的读数普遍没有台账**，所以它在所有管理体系里都是不存在的。

场景四：明星分析师跳槽。答案：明显低于一，而且有一个决定性的补充。格罗伊斯伯格、李与南达对华尔街证券分析师的研究发现，明星分析师跳槽之后表现立即下滑，并且这种下滑持续至少五年；下滑幅度在独自跳槽、跳到实力较弱的平台时最大。而这项研究里最锋利的一个结果是：**带着原来的团队整体迁移的明星，表现没有显著下滑**。记录在哪里：分析师排名是一份公开的、逐年发布的正式台账，这是少数几个 s 可以被真实计算的场域之一。

场景五：用生成式工具代做作品集。答案：接近一。一份可陈述的作品——按规范写就的代码、按模板做成的分析报告、按套路完成的设计稿——被“照单复现”的成本已经趋近于零，留种率接近一。**于是按本文的判断，它必须正在失去溢价。**这是一个可以被数据检验的推论，第十一节报告这项检验。

把五个场景排成一列，可以看到一件事： **s 的高低与“这件事重不重要”完全无关**，与“它写得细不细”也无关。核查表极端重要且写得极细， s 接近一；菜谱写得同样细， s 停在中间；分析师排名是最正规的台账之一， s 明显小于一。**决定 s 的是这项优势的载体结构：它长在构件上，还是长在配组上。**

九、亲本签名

现配优势有一个反直觉的推论，本文称之为**亲本签名**，因为它直接来自育种学的一个基本事实：配出高产杂交种的亲本自交系，各自单独种植时表现平庸，甚至不如普通品种。

推论是：**在一个高度依赖配组的高绩效单元里，其成员的单项能力测评不应当系统性地高于同行基线，甚至可能低于基线。**

这与资源基础观的预测正相反。按资源基础观，持续竞争优势来自有价值、稀缺、难以模仿、不可替代的资源；那么对一个持续高绩效的单元做资源审计，应当能找到某一项或某几项显著优于同行的资源。而按本文的判断，如果优势的载体是配组，逐项审计的结果应当是**每一项都平庸**——优势正好落在审计的格子之间。

于是同一份测评表，两个理论给出相反的读法：一份“各项能力测评均处于同行中位、无一项突出”的报告，资源基础观读作“此单元无持续竞争优势，其当前绩效来自运气或市场条件”，本文读作“这恰恰是配组型优势的签名，应当检查该单元的配组稳定性而不是补齐单项短板”。

这个分歧是可以被观察裁决的，而且裁决的方式是现成的：找一批持续高绩效的小单元，做逐项能力审计，同时记录其配组稳定性（成员组合的变动频率）。如果单项审计普遍平庸而配组稳定性显著高于对照，本文对；如果单项审计能稳定找出突出项，本文错。

亲本签名还有一个直接的实践含义，它与主流的人才管理相冲突：**当一个高绩效小组的成员做单项测评时露出短板，主流反应是补短板——培训、轮岗、换人。**如果这个短板恰是亲本签名，那么补短板的操作会先破坏配组（培训与轮岗都要把人从配组里抽出来），再用一个单项更强但未与该配组磨合过的人替换他。这个操作在所有单项指标上都会显示为改善，而在总产出上显示为下降；而由于失效不产生信号（第七节签名二），下降会被归因于市场、周期或别的什么。

十、与既有理论的分界

本节把本文放到它必须面对的邻居中间，逐条给出**判定形式**：在何种观察下本文对而它错，在何种观察下反之。凡不能写成这种形式的分歧，本文视为措辞之争，不予保留。

一、贝克尔的专用人力资本（1964）。这是离本文最近、也最危险的一位。贝克尔区分通用人力资本与企业专用人力资本，后者只在特定雇主处有价值，因而由雇主与雇员共担投资成本。赫克曼与皮萨诺的观察通常正是被归入这一框架的。

分歧在于**存量与制作**。贝克尔框架下，专用人力资本是一份积累起来的存量，它属于“这个人在这家企业”这一匹配，会随时间累积，也会随离职折旧。本文主张它不是存量而是每次的产物。判定形式：**中断检验**。取一对长期稳定配合的搭档，令其分开一段可观的时间（例如六个月，期间双方均无该配组下的交付），然后复合，观察复合后前若干次交付的读数。存量论预测：优势基本保留，因为资本仍在账上，只有轻微折旧。本文预测：优势明显低于中断前，并在其后按次数（而非按时间）回升。更简洁的一个版本：取两份履历，一份“生涯累计量高、近期在该配组内的量为零”，另一份反之。贝克尔框架判前者更强，本文判后者更强。外科文献中“近期量的解释力强于生涯累计量”的既有方向站在本文一侧，但这尚不足以裁决，因为它可以被解释为折旧速度较快。中断检验能裁决，因为**折旧是时间的函数，重置是次数的函数**：中断后按交付次数回升而非按休整时间回升，只有本文能解释。

二、波兰尼的默会知识（1966）。默会知识写不下来，这一点与格④相同。分歧在于**载体**：默会知识长在人身上，随人迁移；现配优势长在配组里，不随人迁移。判定形式：**换东家检验**。取一位在原处被公认为强的人，让他单独（不带任何原搭档、原客户）迁入一个同类新场，观察其读数。默会知识论预测读数基本保留（手艺在他身上）；本文预测明显下滑并按次数重建。格罗伊斯伯格等人的“单独跳槽下滑、带队跳槽不下滑”的对比，正是这一判定形式的一次真实实施，结果站在本文一侧。但本文不主张默会知识不存在——第六节格③承认它；本文主张的是，在**可陈述层被廉价供给之后**，剩余竞争差中格④的份额高于格③，而这一份额比较是可以同一套读数测的。

三、斯彭斯的信号理论（1973）。信号理论解释的是“为什么可陈述的凭证会贬值”：当获取凭证的成本对高低能力者的差异缩小时，凭证的分离能力下降。这与本文第七节的满单件机制高度相容，本文不与它争这一段。分歧在于**它之后是什么**。信号理论的框架里，能力本身仍是一个先在的、有待被信号揭示的私人信息；本文主张在格④里没有待揭示的私人信息，因为那份优势在交手之前并不存在于任何一方手里。判定形式：如果能设计出一种更昂贵、更难伪造的信号，从而稳定地把强者分出来，则信号理论足够、本文多余；本文预测任何一种可陈述信号在被制度化之后都会在可预期的时间内失去分离力（第十二节的次序条款把这一预测写成了可证伪的形式）。

四、资源基础观（Lippman 与 Rumelt 1982；Barney 1991）。因果模糊性这一概念与本文最为接近：竞争者无法确定优势的来源，因而无法模仿。分歧有两处。其一是**模糊的位置**：资源基础观里模糊性是模仿者的认识论困难（局外人看不清），本文里它是本体论事实（局内人也交不出，因为那一手在交手当场才发生）。其二是第九节的**亲本签名**：资源基础观预测资源审计能找到突出项，本文预测各项平庸。第二处是可观察的裁决点。

五、格罗伊斯伯格等人的可携带性研究（2008）及其后续。本文与它在方向上一致，欠它一份重要的证据。分歧在于**它停在哪里**：可携带性研究记录了“跌多少”，并把带队不跌解释为团队特定人力资本；它没有给出一个可以在任意场域清点的读数，也没有把“照单交接”作为一种独立于人员迁移的检验手段。本文的增量是：把**人员迁移与清单交接**分开。带队跳槽是把整个配组搬走（因此不跌），清单交接是把构件与说明书交出而配组留在原处（因此跌）。这两种操作在可携带性研究里没有被分开，而它们的分离恰是本文命题的核心检验（第十二节 F2）。

六、工人—企业固定效应分解（Abowd、Kramarz 与 Margolis 1999 一脉）。这套方法把工资分解为工人效应、企业效应与残差，是关于“优势归谁”的最严格的经验工具。

分歧极其干脆：这一脉把**匹配效应作为残差处理**（因为在通常的样本中它难以与噪声分离），而本文主张匹配效应就是主项。判定形式：在流动率足够高、能识别匹配效应的样本里，若匹配效应的方差份额随可陈述技能的普及而**上升**，本文对；若其份额稳定或下降，而工人效应份额上升，本文错。这是一个可以用现有行政数据直接做的检验，本文未做，登记为欠账。

七、个人商誉的估值实务。在美国税务法院一九九八年的一则判决中，法院认定与创始人个人相联系的商誉在他未与公司签订不竞争协议时属于其个人资产，而非公司资产。这是实务界对“优势归属”最较真的一个场域。分歧在于**商誉被当作可估值出售的存量**：估值师给它一个数，买方付掉，交割完成。本文预测：这类交易之后，若卖方本人不继续在场，那份被买走的东西会在可观察的时间内衰减；买方所购之物的实际交付需要卖方**逐次到场续制**。判定形式：比较“创始人留任”与“创始人离场”两类交易在交割后若干年的客户留存与收入曲线；存量论预测差异有限，本文预测差异显著且不因交接文档的完备程度而缩小。

八、判断力叙事（Agrawal、Gans 与 Goldfarb 2018）。这一脉的论点是：当预测变得廉价，判断就变得值钱。本文与它的分歧不在方向而在**命名的后果**：一旦“判断力”被命名为一项可培养、可评估的个人官能，它就会被课程化、被评估化，从而落入格②满单件，重演它本要解决的问题。判定形式：观察“判断力”类课程与认证的普及曲线与其薪酬溢价曲线的先后次序（第十二节 F3）。若溢价在普及之后仍能维持，本文错。

九、技能榜单（世界经济论坛及各类就业力框架）。这不是理论而是实务界自己的答案，但它是本文最直接对手，因为它是当前教育政策与课程改革的主要依据。分歧已在第七节陈述：榜单把“重要”与“能区分人”混为一谈。判定形式：取历年榜单上的能力条目，追踪其在职位描述中的出现率与其边际薪酬关联。本文预测两条曲线的次序是——出现率上升在前，边际关联下降在后，最终该条目退出职位描述（因为已成为默认）。若出现率上升伴随边际关联同步上升并维持，本文错。

第二层：这些邻居共有的地基。上述九位彼此激烈争论，但它们全部接受**单一载体假定**：优势必须由某一个单一载体持有——一个人（贝克尔、波兰尼）、匹配（专用人力资本、固定效应）、信号（斯彭斯）、资源（资源基础观）、资产（个人商誉）、官能（判断力）、条目（榜单）。承认“优势可以没有常驻载体、每次由临时配组现做”，这些框架的核算单位就全都失去对象。这就是本文所占的位置：不是在这些框架之间选边，而是把它们脚下那块地翻过来。

十一、三处反向约束

一个判断如果没有被自己的材料削弱过，它就还没有被检验过。本节报告三处反向材料，它们中的两处实质性地改变了本文原本要写的主张。

反向一：可清单化的东西能救命，而且必须被普及。世界卫生组织手术安全核查表是本文材料中最强的反证。它可清单化、留种率接近一、转移成功，并且效果是并发症与死亡率的显著下降。按本文原本要写的强句——“凡可写下者皆无竞争价值”——这份核查表应当是一件无关紧要的东西，而事实是它救了很多人的命。

这一处删掉了那个强句。修正后的范围是：**可清单化的部分统治的是地板，现配优势统治的是地板之上的竞争差**。两者不是价值高低之别，而是功能之别。在安全关键的领域，把格③默会手艺与格④现配区尽可能地压进格②（即让好的做法变成人人照做的清单），**是善不是病**——那正是核查表所做的事。本文的批评只针对一件事：在这些领域之外，把清单化当作竞争力政策的全部，并用清单的长度衡量教育的成效。

这一修正同时改动了本文的价值排序：本文原本把格②写成纯粹的失效格，现在它是**地板的生产格**，只是不能同时充当竞争差的来源。

反向二：上游可以被整体挖走。一九九四年，美国第八巡回上诉法院维持了一项判决，认定被告公司挪用了原告的亲本自交系用于自己的杂交种生产，赔偿约四千六百七十万美元。这个案子对本文有直接的杀伤力：如果优势长在配组里而不长在任何构件里，怎么会有人靠偷构件成功？

答案是：他偷的不是显露的产物（杂交种子），而是**上游的亲本组合**。这一处删掉了本文原本要写的第三条性质的强形式——“**现配优势不可被窃**”。修正后是：**现配优势不可经显露产物转移，但其上游配组载体可以被整体取走。**

这个修正让本文与格罗伊斯伯格的发现严丝合缝地对上了：带着整个团队跳槽的明星表现不下滑，与偷走亲本自交系的种子能配出优势，是同一条通道。它同时给出一条实践含义：对格④优势的真正威胁不是文档泄露、不是作品公开，而是**整块配组被端走**——核心搭档一起被挖、客户关系整体转移、内部工装连人带物被复制。防守的对象因此完全不同：不是保密，是配组的完整性。

反向三：不可预测性有时效。本文的一个隐含前提是配组结果难以事前预测——特殊配合力必须真配一次才知道。但育种学自己正在削弱这一点：基因组选择方法已经能在一定程度上预测杂交组合的表现，从而减少必须实测的组合数。相应地，若组织行为与人事匹配的预测方法有类似进展，一部分现配优势会被移入可预先计算的范围，从而失去其稀缺性。

这一处不改主张，但给主张加了时效条件：**现配优势的不可预测性是一个历史读数，不是一个永恒屏障**。本文的命题应当被理解为对当前条件的描述，其有效期取决于配组预测技术的进展速度，而不是取决于任何原理性的不可能。

三处反向材料的共同效果是把本文从一个关于“什么有价值”的宣言，收窄成一个关于“竞争差的载体结构在特定条件下是什么”的有边界的判断。收窄之后它更弱，也更可能是对的。

十二、一次事先登记的检验

本文的第一层与第二层主张（载体结构、留种率读数）需要长期的实地实验才能裁决，本文在下一节把它们写成可证伪的条款。本节报告一项可以立即用公开数据完成的第三层推论检验。

推论。若第八节场景五成立——可陈述构件的复制成本趋零使其留种率接近一——则可陈述构件的边际报酬应当被压缩。

检验对象。美国全职软件开发者的公开薪酬调查（同一机构逐年发布的开发者调查），比较大语言模型广泛可用之前的二〇一九年与之后的二〇二四年。自变量为受访者所报“正在使用的编程语言条目数”，因变量为年报酬的自然对数，控制专业从业年限、年限平方与本科及以上学历。

事先登记。下列全部条款在读取任何数据之前写入一份文件并计算其散列值存档（SHA-256：
7a336bac3be1fb06112f0a1d137519624abce75dd0e92217d6e84a41f910dac9），此后未作修改：

1. **判负条款：**若二〇二四年的语言条目数系数不低于二〇一九年系数的百分之七十五（且二〇一九年系数为正），则本推论判负。

2. **无效条款**：若二〇一九年系数不为正，则记"跑不动——基线无溢价"，不得改写为有利结论。
3. **不利结果处置**：无论结果如何照原样报告，不得增删控制变量、不得更换年份、不得改变样本筛选。
4. **停止规则**：每年只跑一次回归；样本筛选条件（国别、全职、报酬区间、年限编码）预先写死；数据源依次尝试官方发布页与公开镜像，任一年不可得则记"跑不动——数据不可得"。

结果。两年样本分别为一二七八四人与四三五九人（筛选后）。控制变量如上，最小二乘估计给出：

年份	每多一项语言条目对数报酬的边际系数	标准误	t 值
2019	+0.0024	0.0019	1.3
2024	-0.0050	0.0024	-2.1

即：二〇一九年每多掌握一种语言，报酬高约百分之零点二四，统计上与零无异；二〇二四年每多掌握一种语言，报酬低约百分之零点五，统计显著。按事先写死的判负条款，本推论获支持。

三条必须同时报告的不利侧写。

其一，二〇一九年的基线溢价本身就与零无异。因此本文不得说"技能条目的溢价塌了"，只能说"它本就近乎不存在，而如今转为显著为负"。这一降格是实质性的：原本预期的"由正转零"没有出现，出现的是"由零转负"，后者的机制可能与本文的解释无关（例如广泛涉猎多种语言的人本就更可能处在报酬较低的岗位类型上，而这种构成关系在两年之间发生了变化）。

其二，两年的样本量与人群构成差异很大（后一年样本仅为前一年的三分之一），本文未对构成变化做任何控制。这足以单独解释符号翻转。

其三，"语言条目数"只是可陈述构件的一个粗代理，它同时携带着"广度—深度权衡"的含义，与本文所要测的东西并不同一。

因此这项检验的正确地位是：**一个与本文推论方向一致、但远不足以支持本文的旁证。**它的主要价值不在结论而在形式：判负条款在看到数据之前就已写死，因此读者可以确知本文没有在结果出来之后调整口径。本文的核心主张仍然等待第十三节所列的实地检验。

十三、证伪条款与赌注

以下条款使用第五节清点手册的唯一口径，不另设变体。

F1（核心，交接实验）。取同质任务族，招募六组以上，每组含一名原持有者与两名同资历接收者。A 组接收者仅获得完备交接清单；B 组接收者与原持有者共事两周后独立执行（不获得额外清单）。控制文档质量与任务难度。**本文预测 B 组读数显著高于 A 组，且 A 组读数不随清单完备度提升而趋近于原持有者。**若 A、B 两组无显著差异，或 A 组读数随清单完备度单调趋近于一，本文核心为伪。

F2（中断—重建）。如第十节判定形式一。若中断后的恢复速度与休整时长相关而与交付次数无关，本文为伪。

F3（次序）。对任一被榜单化的能力条目，本文预测三个事件的时间次序为：课程化与认证化的普及 → 边际薪酬关联下降 → 退出职位描述（成为默认而不再列出）。若观察到普及之后关联仍长期维持或上升，本文为伪。

F4（亲本签名）。如第九节。若持续高绩效小单元的成员单项能力审计能稳定找出突出项，且突出项与绩效的关联强于配组稳定性与绩效的关联，本文为伪。

F5（个人商誉交割）。如第十节判定形式七。若创始人离场型交易与留任型交易在交割后的客户留存曲线无显著差异，本文为伪。

F6（杀死条款）。若在六组以上同质案例中，完备清单交接之后接收方能稳定复现原持有者读数的九成以上，并维持一年，则本文核心主张为伪，无须任何补救性修订。

赌注。至二〇三〇年十二月三十一日，在同口径的美国开发者报酬调查中，“可列举技能条目数”的边际报酬关联不会回升为显著为正。若该调查停办或口径不再可比，本赌注记为不成立，不得改判为命中。

尚未清偿的欠账。第十节判定形式六（匹配效应方差份额随可陈述技能普及而变化的检验）可以用现有的行政雇佣数据完成，本文未做。这是本文目前最应当补上的一项。

十四、对教育与雇用的含义

如果上述判断成立，有三组直接后果。

对毕业生。该攒的不是更长的清单，而是**留种率低的经历**：与具体的人、在具体的场子里、围绕具体的问题反复配过，并且当场做出过别人照单复现不了的东西。判断一段经历是否属于这一类，有一个可自查的问题：**如果把我在这段经历里做的事写成一份完备的交接清单交给同班同学，他照做能做到我的几成？**能做到九成的那些经历，写在简历上是必要的（它们是地板），但不要指望它们区分你。

这里有一个必须直说的代价。现配优势不随人整体携带，这意味着**它每换一个场子都要重新长一次**。因此以现配优势为主要资本的人，其职业路径的形状与以证书为主要资本的人完全不同：前者换东家的成本高得多，其价值在长期稳定的配组中复利增长，在频繁跳动中被反复清零。这不是一个励志的结论，它是这个载体结构的直接后果。

对大学。课程可以生产格①与格②，很难直接生产格④，因为格④按定义不可清单化。但大学能做一件关键的事：**为配组提供密度与时间**。真正长出格④的场所是长期项目、稳定的师生配对、连续多学期的合作团队、有真实交付对象的工作——凡是把同一批人反复放在同一个真问题上的安排。反过来，一切按学期打散重组、按模块独立评分、按个人产出计分的设计，都在系统地摧毁配组，即使每一门课的教学质量都在提高。

这里有一个直接的评价体系冲突：格④的产物**无法归到单个学生名下**，而现行的评价体系必须把每一分归到一个人头上。这不是技术难题，是核算单位与所要培养之物之间的结构性错配。本文不提供解决方案，只指出它的位置。

对雇主。三条。第一，招聘时对“清单全勾”的候选人保持警惕：满单件的三个签名（第七节）同样适用于个人。第二，评估高绩效小单元时，**先看配组稳定性再看成员测评**；对成员做单项补短板之前，先估算这次操作对配组的破坏。第三，防守的对象不是文档而是配组完整性：整块团队被挖走、客户关系整体转移，才是真正的失守，而这两件事在现有的保密制度里几乎不受管辖。

十五、反噬、自反与限制

反噬预言。如果留种率进入考核，最省事的达标办法有两条，而且两条都会摧毁它本要保护的东​​西。其一是**表演性配组**：为了让读数好看而制造搭档署名、轮岗打卡、“协作时长”记录，使配组变成一种可申报的形式。其二是**做低清单件的成绩**：既然 s 是接收方读数与原读数之比，那么把清单写得含糊、把交接做得敷衍，就能让 s 变低而显得“我的工作不可替代”。第二条尤其危险，因为它直接奖励知识囤积，而知识囤积正是组织最需要防的事。

本文看不到这个出口。任何一个被用于分配资源的读数都会被优化，而 s 的两个端点（原读数与接收方读数）都在被考核者的影响之下。因此本文的建议是： **s 不应当被用于对个人的奖惩**，只应当被用于诊断——判断某项优势属于哪一格，从而决定是把它清单化（若属地板）还是保护配组（若属竞争差）。这个建议会不会被遵守，本文不乐观。

自反。本文自身是一份可清单化的论述：它有定义、有读数、有清点手册、有判据、有条款。按本文自己的分类，它落在格②的形式里。这意味着：**把本文的骨架抄下来去教学生，其留种率会很低**——不是因为本文写得不好，而是因为读一份关于配组的清单，与在真实配组里做过一次，是两件不同的事。本文预言它自己作为教学材料的效果将显著低于它作为诊断工具的效果。写出这篇东西的那一手配组——那些具体的材料、具体的反例、具体的争论对象在特定时刻被摆在一起的方式——不在文内。

伦理限制。三条，本文视之为使用本读数的前提条件。

第一， **s 指位置不指人**。它测的是“这项优势有多大比例可经清单转移”，不是“这个人有多不可替代”。把 s 读成个人价值的排序是误用，而这种误用几乎必然发生，因此凡使用 s 的组织应当在制度文本中写明这一点。

第二，**不得为了做数而制造人事变动**。测 s 需要交接，而交接意味着换人。在安全关键的系统里（医疗、航空、电力），为测量目的制造轮换是不可接受的；在这些领域，如第十一节所述，正确的方向本来也是尽可能把优势清单化。

第三，**若 s 进入任何影响个人待遇的决定，其编码规则必须公开，并提供复核通道**。一个由第三方判定“清单是否完备”、由既有台账提供读数的程序，其每一步都可能出错；被这些判断影响的人有权看到规则并提出异议。

十六、现配优势在毕业生身上的四种具体形态

前面几节的论证都在结构层面。本节把它落到一名二十二岁的毕业生身上：格④在这个年纪、这种处境里，具体长成什么样子？以下四种形态不是穷举，而是四个可以被识别、可以被有意经营、并且都能用留种率自查的位置。

形态一：与一个具体的人配熟。这是最基本的形态，也是最容易被简历格式抹掉的一种。一名学生跟同一位老师做了三个学期的课题，或者与同一位同学从大二一起做到毕业，或者在实习期间固定与同一位工程师结对——这种关系里长出来的东西，在简历上只能写成“参与某项目”，而它的实质是：**两个人之间形成了一套只对这两个人有效的省略**。该说什么不说什么、哪一步可以跳过、对方的判断在哪里可靠在哪里需要复核、出问题时先看哪里——这些都不在任何交接文档里，因为它们不是任何一方的知识，是两人之间的约定，而且大部分是默默约定的。

自查：把你与这位搭档的合作方式写成一份交接说明，交给一位同水平的同学去接替你的位置，头三次交付他能做到几成？如果接近十成，说明你们之间的所谓默契其实是流程，那是格①；如果明显做不到，且做不到的部分你自己也说不清在哪里，那是格④。

形态二：与一个具体的场配熟。一家公司的代码库、一个实验室的仪器脾气、一个客户的历史遗留问题、一所学校的行政路径——每一个真实的场都带着大量不会被写进文档的具

体性。赫克曼与皮萨诺所观察到的"随该院手术量而非总量改善",说的正是这件事:外科医生与那家医院的护士、器械、流程、病人构成之间形成的配合,不能带走。

对毕业生的直接含义是:**在一个场里待到"我知道这里的东西为什么是这样"的程度,与在五个场里各待三个月,产生的是两类不同的东西。**后者产出的是可陈述的见闻(格①),前者产出的是格④。这不是说不该换,而是说换的代价必须被明码算进去:每换一个场,这部分资产从零开始。

形态三:一手把便宜构件配起来的组装法。这是最贴近人工智能时代的一种形态。当所有工具、模型、模板、代码片段都廉价可得时,构件本身完全不构成优势——每个人都有同样的积木。差别出现在:面对一个具体的、脏的、边界不清的真问题时,谁能判断该用哪几块、按什么顺序、在哪里必须自己动手不能交给工具、在哪里工具的输出必须被复核。

这一手极难写下来,原因不是它神秘,而是它**依赖于问题的具体性**:同样的组装法换一个问题就是错的,所以任何写下来的版本要么太具体(只对那个问题成立)要么太抽象("要具体问题具体分析",等于没说)。这正是可清单化失败的典型机制。

自查特别简单:把你解决某个真问题的完整过程写成教程发出去,看别人照做能不能在新的类似问题上做成。教程被大量点赞而后续没有人真的靠它做成事,通常说明你真正做的事没写进去。

形态四:一段被外部承认的连续交付史。前三种形态都有一个共同弱点:它们没有台账(第八节场景三)。第四种形态是前三种的可见化——不是把优势本身写下来(写不下来),而是让配组的**产物序列**被外部看见:连续维护了两年的开源项目、为同一个真实用户群持续迭代的作品、一个有真实读者的长期写作、一份被同一个客户反复续约的服务。

这一形态的价值不在于它证明了你会什么,而在于它是格④唯一能被外部核验的形式:**别人无法核验你的配组,但可以核验这个配组持续产出过东西。**这也是本文对"作品集"的一个修正建议:作品集的关键不是作品的数量与精美度(那已经可以廉价生成),而是**同一条线上的连续性与真实对象**。一份由十个互不相关的漂亮成品组成的作品集,与一份由同一个项目两年间十次迭代组成的记录,在本文的框架里是两个格的东西。

四种形态的共同点是:它们都需要**时间与重复**,而且都无法通过增加课程、认证或工具来加速。这是一个不受欢迎的结论,因为它意味着这部分竞争力不能被购买、不能被速成、也不能被并行地大量获取。它也解释了一个常见的困惑:为什么很多"简历很满"的毕业生在工作中表现平平,而某些"简历很薄"的人很快就显出差别——前者把四年花在了扩充格①与格②上,后者把四年花在了两三个配组里。

十七、四个可能的误读

本文的主张容易被读成四件本文并未主张的事,逐条澄清。

误读一:"所以证书和技能没用。"本文不主张这一点。格①是地板,而地板决定的是你能否进入某个场——没有基本技能,你连参与配组的资格都没有。本文的主张是关于**地板之上的**:在满足门槛之后,继续加长清单的边际回报会迅速趋近于零,而同样的时间投入配组会有明显更高的回报。这是一个关于资源分配的判断,不是一个关于价值的判断。第十一节的核查表案例已经说明,在安全关键领域,清单化甚至是首要的善。

误读二:"所以竞争力就是人脉、关系、圈子。"本文与这一读法的差别是决定性的。人脉是一份**关系的存量**:认识谁、能找到谁、在谁那里说得上话——它可以被清点,可以被继承(介绍信、引荐),在一定程度上也随人迁移。现配优势不是关系本身,而是**关系被使用时产生的东西**:同样两个人在一起,可以只是认识,也可以在一次次真实的交付中长出那套

只对他们有效的省略。前者是名单，后者是产物。一个人可以有极长的名单而完全没有格④，这在实习经历丰富而每段都很短的毕业生身上很常见。

误读三："所以不必努力，运气而已。"恰好相反。运气不可复现，而现配优势在同一配组内高度可复现——这正是它能被称为优势的原因。它与运气的分界可以用一句话划出：**换一个搭档、换一个场，运气的期望值不变，现配优势归零。**归零这件事本身就说明它不是随机的：如果它是随机的，换一个场它应该同样有一半机会出现。现配优势要求的努力量并不比攒清单少，只是努力的形式不同——它要求在同一处待得足够久、配得足够深，而这在一个鼓励广度与流动的环境里恰恰是更难的选择。

误读四："这只是把'团队合作'重新说了一遍。""团队合作"是一项被课程化、被评估化的能力，它的载体是个人（一个人"善于团队合作"），它可以写进简历，可以被培训，可以被打分——按本文的分类，它是标准的格②满单件。本文所说的现配优势不是一项个人能力，它没有个人载体：它不在你身上，也不在你搭档身上，它在你们之间，并且只在你们真正交付东西的时候存在。一个"团队合作能力评分很高"的人，换一个团队仍然评分很高（因为那是他的属性）；一个拥有格④优势的人，换一个团队要从头长。这两件事在观察上是可以分开的，第十节判定形式二给出了分开的方法。

最后一条限制。本文的全部论证依赖一个经验前提：可陈述认知劳动的复制成本正在趋近于零。这个前提在当下看起来牢固，但它是一个关于技术与市场的判断，不是一条原理。如果这一前提在某个方向上逆转——例如可陈述产出的核成本急剧上升，使得"照单做出来的东西"重新变得不可信任因而稀缺——那么格①会重新获得区分力，本文的主张的适用范围会相应收缩。本文把这一条写在这里，是为了让读者知道：**本文不是关于人的本性的论断，而是关于当前一组条件下竞争差落在何处的论断。**条件变了，结论应当跟着变。

十八、结语

本文的主张可以压缩成一句话：**在可陈述的东西变得极其便宜之后，人与人之间剩下的差别，不是谁攒得多，而是谁能在当场把便宜的东西配成不便宜的结果。**

这个判断的代价是它取消了一件很多人依赖的东西——那份可以随身携带、可以写在纸上、可以在任何地方兑现的个人存量。本文认为那份存量在很大程度上是一个记账习惯的产物：我们把一次次配组的产物记在了其中一个构件的名下，久而久之就以为它长在那里。杂交稻的产量记在种子袋上，转让费记在盘店的合同里，双打的成绩记在两个人各自的排名里——记账总要有个名下，而名下久了就成了归属。

如果这个判断是对的，那么对一名毕业生最诚实的建议是一句不太好听的话：你真正值钱的那部分，将不会出现在你的任何一份正式记录里；它只会在整体录用你的那一刻，被当作一笔说不清明细的钱付掉——而且下一个场子里，它还要重新长一遍。

附录 留种率的一次示例清点

为使第五节的清点手册可以被别人照着做一遍，本附录给出一个完整的示例。示例中的情境是构造的，但每一步的判定规则都取自手册原文，不作简化。

情境。某软件公司的一个四人小组，两年来负责同一家企业客户的数据接口维护。组内一名工程师甲被公认为该客户事务上最关键的人：客户方的历次变更请求由他判定可行性，出现故障时由他定位，客户方的技术负责人只找他。甲将于三个月后调岗。公司想知道：甲身上这部分优势，有多少能经交接留下来？

第一步，确定同类任务与既有读数。按手册第一、第四条，同类任务限定为"该客户的变更请求处理"，不含新客户接入、不含内部重构。读数取组织本来就在记录的两项：变更请求

的一次通过率（客户验收无返工的比例）与从受理到交付的中位时长。这两项均在工单系统中已有两年记录，不为本次测量新造。甲最近半年的读数为：一次通过率 0.82，中位时长 3.5 个工作日。

第二步，判定资历同一档。按手册第二条，接收方须与甲在学历、专业年限、职级与相关认证四个公开维度上同档。公司找到的两位候选人中，一位年限少两年（不同档，排除），另一位同档，记为乙。

第三步，完备清单与其判定。甲用四周时间写出交接材料：接口文档、历史变更清单与判定依据、常见故障的定位路径、客户方各联系人的职责与偏好、监控与告警的解读说明、可复用的脚本与模板。按手册第三条，完备性不由甲自评，而由不参与该客户事务的第三名工程师按预先约定的条目框架逐项检查，结果为缺项一项（客户方财务审批流程的特殊情形未覆盖），缺项率低于一成，判定为完备清单，测量有效。

第四步，执行与读数。乙独立处理该客户的前三次变更请求，不得向甲求助（这一条必须写进方案，否则测的是共事而非清单）。三次的读数为：一次通过率 0.33（三次中一次无返工），中位时长 8 个工作日。

第五步，计算。

指标	甲（原持有者）	乙（清单接收方）	比值
一次通过率	0.82	0.33	0.40
中位时长（越小越好，取倒数比）	3.5 日	8.0 日	0.44

两项读数给出的留种率分别为 0.40 与 0.44。按手册，两项均报告，不合成单一分数——合成会掩盖两项指标可能分歧的情形。本例中两项一致，可以说：**该项优势经完备清单交接后，接收方复现了约四成。**

第六步，读出结论。s 约为 0.4，落在第六节的低留种率一侧；而该优势的可清单化程度经第三方判定为高（缺项率低于一成）。二者结合，落在格④而非格②——**这不是文档不足的问题，是载体不在文档里的问题。**

于是管理动作的方向就与直觉相反。直觉的动作是“交接材料还不够细，让甲再补”；本例的读数说明再补的边际收益很低（清单已判完备而 s 仍为 0.4）。按本文的框架，正确的动作有两个方向：其一，若该客户事务属于地板（不出错比出彩重要），则应当投入把不可清单的部分设法清单化——即第十一节所说的核查表方向，代价是长期投入，收益是稳定；其二，若该客户事务属于竞争差（这家客户之所以留下来正因为这份处理质量），则应当**保护配组**而非加厚文档：推迟甲的调岗、让乙与甲并行共事而非顺序交接、或者接受这份优势将在乙手里重新长一遍并为此预留时间与容错。

这个示例暴露的三处困难，本文如实登记。

其一，“前三次”可能太少。三次读数的抽样误差很大，尤其在一次通过率这类二值指标上（本例中 0.33 实际上是三分之一）。手册规定取前三次是为了容纳适应成本，但这与统计精度直接冲突。改进方向是在任务频率允许时取前十次并报告置信区间；本文未给出这一改良的正式规定，登记为欠账。

其二，“不得求助”这一条在真实组织里很难执行，而一旦允许求助，测到的就不是清单的留种率而是清单加咨询的留种率。这两者都有意义，但必须分开报告。

其三，比值形式对“越小越好”的指标需要取倒数，而对没有自然零点的指标（如满意度评分）无法直接取比值。本文目前只对有自然零点的比率型与时长型指标给出计算规则；其他指标类型的处理规则尚未写定。

以上三处都不是本文核心主张的漏洞，而是读数工具的粗糙之处。它们被写在这里，是为了让任何想复现这套清点的人事先知道自己会撞上什么。

参考文献

Abowd, J. M., Kramarz, F., & Margolis, D. N. (1999). High wage workers and high wage firms. *Econometrica*, 67(2), 251–333.

Agrawal, A., Gans, J., & Goldfarb, A. (2018). *Prediction Machines: The Simple Economics of Artificial Intelligence*. Harvard Business Review Press.

Barney, J. (1991). Firm resources and sustained competitive advantage. *Journal of Management*, 17(1), 99–120.

Becker, G. S. (1964). *Human Capital: A Theoretical and Empirical Analysis, with Special Reference to Education*. Columbia University Press.

Codman, E. A. (1914). The product of a hospital. *Surgery, Gynecology and Obstetrics*, 18, 491–496.

Dranove, D., Kessler, D., McClellan, M., & Satterthwaite, M. (2003). Is more information better? The effects of “report cards” on health care providers. *Journal of Political Economy*, 111(3), 555–588.

Groysberg, B. (2010). *Chasing Stars: The Myth of Talent and the Portability of Performance*. Princeton University Press.

Groysberg, B., Lee, L.-E., & Nanda, A. (2008). Can they take it with them? The portability of star knowledge workers' performance. *Management Science*, 54(7), 1213–1230.

Haynes, A. B., Weiser, T. G., Berry, W. R., et al. (2009). A surgical safety checklist to reduce morbidity and mortality in a global population. *New England Journal of Medicine*, 360(5), 491–499.

Huckman, R. S., & Pisano, G. P. (2006). The firm specificity of individual performance: Evidence from cardiac surgery. *Management Science*, 52(4), 473–488.

Kewanee Oil Co. v. Bicron Corp., 416 U.S. 470 (1974).

Lippman, S. A., & Rumelt, R. P. (1982). Uncertainty imitability: An analysis of interfirm differences in efficiency under competition. *Bell Journal of Economics*, 13(2), 418–438.

Martin Ice Cream Co. v. Commissioner, 110 T.C. 189 (1998).

Pioneer Hi-Bred International v. Holden Foundation Seeds, Inc., 35 F.3d 1226 (8th Cir. 1994).

Polanyi, M. (1966). **The Tacit Dimension**. University of Chicago Press.

Restatement (First) of Torts § 757 (1939), comment b.

Shull, G. H. (1908). The composition of a field of maize. **Journal of Heredity**, 4(1), 296–301.

Spence, M. (1973). Job market signaling. **Quarterly Journal of Economics**, 87(3), 355–374.

Sprague, G. F., & Tatum, L. A. (1942). General versus specific combining ability in single crosses of corn. **Journal of the American Society of Agronomy**, 34(10), 923–932.

Uniform Trade Secrets Act (1979, amended 1985), § 1(4).

World Economic Forum. (2025). **The Future of Jobs Report 2025**. Geneva.

Stack Overflow. Annual Developer Survey public data sets, 2019 and 2024 waves.

第八章 脱稿步

【章首】本章给方程添的第四条本体性质：**续签性**——署名是一份没有到期日、却已经会到期的凭据，效力只能在自己产物的延长线上被当场重做；在人机连续参与的产物上，“谁做的”并不存在一个等待查明的历史事实，针对本人产物的追问因此不是取样，而是该量的**生成**。读数是脱稿成立率 q 。本章的独门证据是三项字段审计：两份互相独立的学习记录标准合计八十四类，无一为“延长线上的即时应答”设过带类型字段，而全部失效机制的触发条件都在签发方一侧——记录体系只有签发方时间，没有持有者时间。本章也带着十章里最坦白的一处失误登记：交互式口头评估这位现役正主是写完后复查才找到的，四条分离（验证与构成之别为首）随失误同页。本章欠的账：判者一致性预实验（功效计算已备，约需一百二十次追问）与用人侧标准的第四次审计，均未做。

大白话：署名要当场重新签一次才算数

一个厨师拿了金奖，奖状挂在墙上。奖状不会过期。但如果今天你走进他的后厨，指着他的招牌菜说：“就着这个，现在给我配一道搭它的菜。”——他能不能当场做出来，是另一回事。

《脱稿步》讲的就是这个差别：**发生留下的署名，没有到期日，却已经会到期**。它的效力不在档案里，只能在你自己产物的延长线上被当场重做一次。

这话听起来抽象，落到操作上却极其简单：**拿着他做的东西，问下一步**。

不是问“你当时怎么做的”（那是回忆，可以背），不是问“你怎么看这个行业”（那是观点，可以准备），而是**指着他交出来的这个东西的具体某处，问它往前一步该怎么走**。续得上，署名有效；续不上，档案再厚也只是一张褪色的收据。

为什么今天要紧？

因为在人和机器共同参与的产物上，“到底是谁做的”这个问题，**没有一个等待被查明的历史事实**。你查不出来，任何工具都查不出来，因为那条界线在过程中就是混的。

这是很多人还没接受的一件事：查重、检测、痕迹分析，全都是在追问一个已经不存在答案的问题。而《脱稿步》给了一条出路：既然“谁做的”查不出来，那就不要查——**当场再让它发生一次**。这时候“谁做的”不再是被查明的，而是被**生成的**。

它不是什么？

它不是面试。面试问的是关于你的问题；脱稿步问的是关于**你交出来的这一件东西**的问题，而且必须是无脚本的、临时的、从具体某处出发的。

它也不是“抽查”或“防作弊”。防作弊的心态会把它做成对抗；它的正确姿态是**验收**——像收房时业主拿着图纸让工人当场演示某个开关一样自然。

读数是什么？

叫**脱稿成立率 q** ：针对本人产物发起的若干次追问中，续签成立的比例。这里有一个书里查出来的、让人有点意外的事实：现有的学习记录标准与人事系统里，**没有一个字段能记下“这个人在自己产物的延长线上当场应答过”这件事**——所有字段记的都是签发方的时间与结论，没有持有者的时间。这不是疏忽，是那个时代不需要它。现在需要了。

一个小场景。

招聘最后一轮，两位候选人交上来的作品集难分高下。面试官不再问"讲讲你这个项目"，而是打开候选人自己的作品，指着第 7 页的一个取舍问："如果这里的数据是反的，这一整页要怎么改？"

甲愣了两秒，说出了三种改法，并指出第二种会破坏第 12 页的一个前提。乙给出了一段流畅、正确、但明显是重述第 7 页已有内容的回答。

差别不在口才。差别在于：甲还站在这个产物的延长线上，乙已经不在了。

再看两个身边的例子。

教孩子做题。孩子交上来一份全对的作业。你想知道他是不是真会，不需要审问，也不需要查他有没有抄——你只要指着其中一道题问："这一步如果换成另一个数，会怎么样？"会不会，三十秒就分明了。**注意你没有去追查过去，你是让它当场再发生一次。**

买二手车。卖家给你一沓完整的保养记录，看起来无懈可击。有经验的买家不会花时间去验证每一张单据的真伪，他会说："发动机盖打开，你现在跟我讲讲这里为什么换过。"能当场讲通、并且讲到你没问的地方，那沓记录才算数。

三种常见的误读。

第一种：以为是在抓作弊。这是最容易走偏的一种，而且一旦走偏，整套设计就会失效。因为对抗的姿态会让对方进入防御，而防御的人会准备——准备好的表演正是这套设计要绕开的东西。正确的姿态是验收：不是"我要证明你不行"，而是"我要看看这个东西还活着没有"。

第二种：以为可以事先给题。给了题就不叫脱稿。这一点没有折中余地：追问一旦题库化，被测的东西当场消失，剩下的只是另一种可以准备的产物。

第三种：以为答不上来就等于没做。不一定。一个诚实的"这里我确实没想过，但我会这样去弄清楚"，往往比一段流利的重述更能说明他在延长线上——因为他知道自己站在哪里，也知道下一步往哪走。真正的失败信号不是"不知道"，是**只能复述已有内容、无法往前走一步。**

给自己做个小检查。打开你最近交出去的一份东西，随便翻到一页，指着任意一处问自己："如果这里的条件变了，往下怎么走？"答得出来的地方，是你真的做过的；答不出来的地方，你心里其实早就知道。

它和旁边几个概念的区别。

和《他手无事账》最容易混，两者也确实经常一起用，但变量完全不同：这一章的变量是**从哪里出发**（必须从本人产物的延长线出发），那一章的变量是**时刻由谁定**（必须不由被读者选）。一次事先约好的、就着本人作品的追问，仍然算脱稿步（因为它从他的产物出发、无脚本），但不算他择（因为时刻是约好的）。四种组合都存在，两轴正交。

和《盲遍》的分工在时间轴的两端：盲遍读形成期，脱稿步读事后。一份档案再厚，如果从没被追问过，那么它的 q 根本没有定义——不是零，是没有定义。

和《底流》有一个有趣的接口：每一次成功的续签，本身就是一次三成分齐备的执行（有真实后果、无辅助、亲手）。所以经常被追问的人，顺带也在被喂。

一句话记住它：这一章管的是**还接不接得上**——你和你交出去的那个东西，现在还连着吗。

怎么长出来：方法、技术与流程

个人层面。

交出任何东西之前，给自己做一次“下一步测试”：**指着自己作品里的任意三处，问自己下一步怎么走。**答不上来的地方，就是你其实没有真正参与的地方。这个测试三分钟就能做完，而且它比任何自查都诚实。

大学发生学教育：把答辩从“讲述”改为“续做”。

这一条是本书里对大学最直接可用的一条，因为它不需要新增任何环节——现有的答辩、汇报、作业面批，全都可以直接改口径。

- **延长线追问，取代叙述性提问。**禁止的问法：“请介绍你的项目”“你为什么选这个题”。要求的问法：“你这里第三步的取舍，如果条件变成 X，怎么走？”“把这个结论往前推一步，会碰到什么？”前者可背，后者不可背。
- **十分钟当场续做。**在提交后随机安排一次十分钟的当场应答，就着提交物本身。这不是加考试，是把已有的答辩时间挪去用。
- **题库禁令。**一旦追问被题库化，这条设计当场自毁。因此规程要写死：**追问必须由评审就着本人产物临场生成，不得有题库、不得提前提供样题。**
- **判者一致性训练。**两名教师就同一次追问独立打分，定期核对一致性。一致性不达标时，问题在评分规则而非学生——这一条能防止“脱稿步”变成教师的主观印象分。
- **持有者时间字段。**在成绩档案里加一栏：本学期被就自己产物追问的次数与成立率。同样地，**不进排名**，只作学生自己的记录。

教师的动作。

第一，追问必须从**学生的东西**出发，不从教师的知识出发——问“你这一页”，不问“教材第几章”。第二，允许学生说“我不知道，但我会这样找”——那也是一种成立的续签，而且往往是最真实的一种。第三，教师自己要示范一次被追问：让学生就教师的讲义临场提问，教师当场续做。

最容易做坏的地方。

变成刁难。追问的目的不是问倒学生，是看他还在不在自己的产物的延长线上。一个善意的、就事论事的追问，和一个想显示自己更高明的追问，在学生那里的效果完全相反：后者会逼出准备好的表演，而准备好的表演正是这一整套设计要绕开的东西。

原文

脱稿步：论人工智能时代大学毕业生核心竞争力的载体转移

——从“可验证的作者身份”到“可续签的署名”，兼对三份记录标准的字段审计

版本：第四版（正式版） 完稿日期：2026年9月 参照语料截止：2026年9月

分类：教育经济学 · 劳动力市场信号 · 教育测量 · 学习记录互操作标准

摘要

研究问题：当交付物可以被廉价代做、可陈述技能可以被普遍供给，大学毕业生凭什么被认作有能力？

核心论点：本文提出，人工智能时代大学毕业生的核心竞争力，不是可陈述技能的加权总分，不是名下产物的存量与光泽，不是落入有利条件场的先在位置，**也不是一次被验证过的作者身份，而是可续签的署名**——署名是一份没有到期日、却已经会到期的凭据，其效力只能在自己产物的延长线上被当场重做。重做的那个动作，本文称为**脱稿步**：在自己名下产物的延长线上、无脚本走出可核的下一步。本文主张，在人机连续参与的产物上，“谁做的”并不存在一个等待被查明的历史事实；因此针对本人产物的当场追问不是对某个既存量的取样，而是该量的**生成**。

方法：（一）理论构造。取三个与劳动力市场无关且彼此不相容的领域——数量遗传学的选择指数理论、攀登伦理中的风格判定、群落生态学的中性理论——识别其共同前提，并以其中第一领域自身的实测结果与制度事实推翻该前提。（二）三项事先登记的字段审计。对两份来源相互独立的学习记录互操作标准（美国主导的一份、欧盟委员会的一份）合计 84 个数据模型类，检验两个问题：是否存在记录该动作的带类型字段；凭据失效的触发条件位于哪一侧。判负条款、不利结果处置与停止规则均在读取数据之前写死。（三）制度场景比较。以一句不含程度词的判据询问五类用人与评价制度。

主要发现：（1）两份标准合计 84 个类中，**无一类**为“在本人既有产物延长线上的即时、非预设应答”设立带类型、可比较、可加权的字段；（2）两份标准**均**设有完整的失效机制族（有效期、凭据状态、审查日期、刷新服务），而其触发条件**全部位于签发方一侧**，无一绑定于持有者当下的表现——本文称之为记录体系“只有签发方时间，没有持有者时间”；（3）五类制度中，仅民航飞行员定期检查一类将持有者时间写入正式字段，其存在依赖于失效后果极端、监管强制、成本全行业分摊三个条件同时成立。

已被自己推翻的两处：第一项审计中“标准里没有任何可容纳该记录的位置”这一强主张，按本文事先写死的条款判负（存在自由文本插槽且数据模型可扩展）；“课堂追问一律不进入任何字段”这一事实判断亦被推翻（交互式口头评估有评分量表并计入课程成绩）。两处均在正文原样保留并给出收窄后的主张。

局限：三项审计均为文本层的类目审计，非属性穷举；判者一致性阈值为设计值而非实测值（第 6.6 节给出预注册的检验方案与功效计算）；核心因果主张尚无个体层面的追踪数据支持。

关键词：可续签的署名；脱稿步；脱稿成立率；代走光泽；签发方时间与持有者时间；学习记录互操作标准；承认与存留

Title: The Off-Script Step: On the Displacement of Graduate Competitiveness in the Age of Artificial Intelligence — From Verifiable Authorship to a Renewable Signature, with Field Audits of Three Record Standards

Abstract

Problem. When deliverables can be cheaply produced on a person's behalf and articulable skills are universally supplied, on what basis is a graduate recognised as competent?

Argument. This paper proposes that the core competitiveness of graduates in the AI era is neither the weighted sum of articulable skills, nor the stock and polish of artifacts bearing one's name, nor a prior position in a favourable environment, **nor a once-verified authorship**, but a **renewable signature**:

authorship is a credential with no expiry date that nevertheless expires, and its force can only be remade, on the spot, on the extension line of one's own artifact. The act of remaking it is the **off-script step** — an immediate, unscripted, checkable next move on one's own prior work. Where human and machine contribution is continuous, no pre-existing fact of authorship awaits discovery; probing anchored to one's own artifact therefore does not sample a quantity, it generates one.

Method. (i) Theory construction from three mutually incompatible source fields — selection-index theory in quantitative genetics, style judgement in climbing ethics, and neutral theory in community ecology — whose shared premise is overturned using empirical results and an institutional fact drawn from the first of them. (ii) Three pre-registered field audits across two independently developed learning-record interoperability standards (84 data-model classes in total), asking whether a typed field exists for such an act, and on which side the triggers for credential invalidation lie. Falsification conditions, handling of unfavourable results, and stopping rules were fixed before data inspection. (iii) A five-scenario institutional comparison driven by a single modality-free diagnostic question.

Findings. (1) Of 84 classes, none provides a typed, comparable, weightable field for an immediate, unscripted response on the holder's own prior artifact. (2) Both standards carry a full family of invalidation mechanisms — expiry date, credential status, review date, refresh service — whose triggers lie **wholly on the issuer's side**, none bound to the holder's present performance. Records keep **issuer-side time** but not **holder-side time**. (3) Of five institutional scenarios, only recurrent airline pilot checking writes holder-side time into a formal field, and does so only where extreme failure consequences, regulatory mandate, and industry-wide cost sharing coincide.

Two claims the paper falsifies against itself. The strong form of finding (1) — that no place exists in the standard to record such an act — fails the paper's own pre-registered condition (free-text slots exist and the model is extensible). The factual claim that classroom probing never enters any field is also refuted (interactive oral assessment is rubric-graded and enters course marks). Both are retained verbatim with narrowed replacements.

Keywords: renewable signature; off-script step; off-script rate; borrowed polish; issuer-side versus holder-side time; learning-record interoperability standards; recognition versus survival

1 引言

1.1 一个已经不再自动成立的推论

两份简历摆在桌上。甲的档案很厚：绩点高、三段实习、一个开源项目、一份数据分析报告，报告干净，图表配色考究，结论也站得住。乙的档案薄得多：绩点中游、一段实习、一个没做完的小工具。按任何一份现行评分表，甲赢，且赢得没有争议。

面试当天，面试官不问通用题，只问甲自己那份报告："你这张图里第三季度的跳点，如果把统计口径换成按下单时间而不是按发货时间，跳点还在不在？在的话，你的第二个结论要改哪一句？"甲说，这个我得回去看一下数据。

同一问题的等价物问乙："你这个小工具卡在哪儿了？现在给你半小时，你先改哪一行？"乙当场说出卡在什么地方、为什么那条路走不通、他试过哪两种绕法、哪一种更省事但会在什么时候咬人。

三年前，这场追问不会发生，因为不必发生：**能做出那份报告的人，几乎必然答得上那个跳点。**本文处理的正是这条推论失效之后的问题。

1.2 失效的不是某一种证据，而是一整族证据的共同挂点

一切大规模筛选制度都建立在一个不曾言明的便利之上：看东西比看人便宜。看一个人要时间、要在场、要反复；看一件东西只要一眼。于是社会发明了一整套用东西代表人的办法——文凭代表受过训练，作品集代表做得出来，绩效记录代表干得成，推荐信代表别人验过。

这套办法之所以曾经可靠，不在于东西本身，而在于一条隐含的约束：**做出那件东西所需要走的路，只能由本人走。**产物因而是路径的一个无偏估计——有噪声，无系统偏差。

本文的第一个论点是：这四类证据并非四种独立证据，**它们共享同一个挂点**；因此当挂点松动时，它们不是逐个失效，而是同批失效。这解释了一件让许多机构困惑的事实：为什么招聘、升学、评审、绩效考核会在同一个时间段内一起出问题。

需要立刻排除一种误读。**这里没有人在造假。**学历是真的，作品是真提交的，绩效是真入账的，所有形式审查都能通过；工具使用可以完全合规、完整披露。变化的只有一件事：做出它所需要走的那条路，缩短了、外包了，或者根本不必走了。因此本文与"学术诚信"或"作弊检测"文献处理的不是同一个对象（详见 2.3 节）。

1.3 本文的贡献

本文提出四项可分别评估、可分别推翻的贡献：

贡献一（本体论层，最强、最易倒）：在人机连续参与的产物上，"谁做的"不存在一个等待被查明的历史事实。署名因此不是可核实的事实，而是**需要被反复重做的凭据**。竞争力的落点随之从"能力存量"移到"署名的现行效力"。

贡献二（测量层）：给出**脱稿成立率 q** 这一读数、七条编码规程与一份可直接使用的记录表（第 4 节、附录 A），使贡献一可被清点。

贡献三（诊断层）：给出以"档案厚薄 $\times q$ 高低"为轴的二维辨别装置，识别出**代走光泽格**——档案很厚而延长线走不出来。该格具备三个签名：短期指标全面改善、失效不产生信号、越正规化越严重（第 5 节）。

贡献四（经验层，最稳、最易检验）：对两份独立标准合计 84 个类的三项事先登记审计，得出"记录体系只有签发方时间，没有持有者时间"这一可复算的结论（第 6 节）。

四项贡献是**分层的**：贡献一若被推翻，贡献四仍然成立，且其本身已是一件值得动手去补的事。

1.4 结构安排

第 2 节回顾邻近文献与现役实务，指出尚未被占据的位置；第 3 节构造理论；第 4 节给出测量；第 5 节给出辨别装置与机制解释；第 6 节报告三项审计与一项预注册的待做检验；第 7 节报告五类制度场景；第 8 节讨论；第 9 节列出三处反过来削弱本文的约束；第 10 节局限；第 11 节证伪条件与一条写死日期的赌注；第 12 节自反性说明；第 13 节结论。

2 文献回顾与本文的定位

2.1 三种现成回答及其共同的失效点

关于"人工智能时代毕业生靠什么"，现有回答大体分三种。

第一种：更高级的技能。低阶、重复、可陈述的工作被接管，剩下分析思维、创造性思维、复杂问题解决等（World Economic Forum, 2025）。其失效方式是内生的：**清单一旦公开，所列各项即开始被课程化、被批量供给**。稀缺不源于高级，而源于难获得；获得路径一旦写成教学大纲，稀缺物即变为入场券。清单越权威，磨平越快。

第二种：与机器协作的能力。其失效方式同样内生：工具厂商的核心目标正是压低使用门槛；且这类能力天然可陈述、可演示、可写入简历，因而立刻落回第一种的命运。

第三种：位置与关系。行业、地域、家庭、时机决定大半（Thurow, 1975; Hirsch, 1976）。这一支多半正确，但回答的是另一个问题：它解释**谁能存留**，不解释**谁被认作有能力**。

三者的共同失效点是同一个：**它们都默认，那个要被评价的东西还能显露出来**。争论集中在"该看哪一样"，无人追问"这几样还看得见吗"。

2.2 十种邻近说法与本文的分界

表 1 逐条给出判定形式。本文的立场是：**"更强调""更深入""视角不同"不构成分界**；说不出判定形式，即等于承认被覆盖。

表 1 与十种邻近说法的分界及其判定形式

#	邻近说法	它说到哪一步	分离线	判决性对照预测
1	专用人力资本 (Becker, 1964)	有些技能只在特定组织值钱	它讲技能的 适用范围 ，本文讲技能与持有者之间的 绑定是否可见	不换组织而档案由工具代做：本文判其落靶格，它预测无变化
2	隐性知识 (Polanyi, 1966)	我们所知多于所能言说	隐性知识是 存量 ，脱稿步是 显露事件	完全可言说的知识亦能产生脱稿步（当场做出未准备的推论），该族预测不出
3	信号理论 (Spence, 1973)	昂贵到不可伪造的信号才可信	它测信号的 成本 ，本文测信号与所指之间的 绑定	昂贵信号成本回升（如恢复闭卷手写）而 q 不回升 → 本文对、它错
4	让分原理 (Zahavi, 1975)	累赘因昂贵而诚实	它要求信号昂贵，本文要求信号 不可预备	便宜而不可预备的追问：本文预测有判别力，它预测没有
5	位置财与岗位竞争 (Hirsch, 1976; Thurow, 1975)	文凭是排队号码牌	它们说排队，本文说 队列已不可读 ；可共存	队列位置分布不变而承认标准转移的时期（本文预测工具普及后 3-7 年）
6	适应性专长 (Hatano & Inagaki, 1986)	惯例型专长 vs 适应型专长	那是对 个体 的分类，本文是对 账本 的诊断	两位适应性专长相同者，机构一设追问一不设：本文预测承认差异巨大，该族预测无差异

7	可托付专业活动 (ten Cate, 2005)	按活动类别授予托付等级	它 前瞻 、按类别；本文 回溯 、锚定此人这一件产物，判据在追问那一刻生成	托付等级高而答不出关于自己上周病历的延长线追问：它判合格，本文判落靶格
8	挑战应答式认证（密码学）	不可重放的现场应答	结构同构；它验证 持有静态密钥 ，本文验证 走过一段不可复制的路	以随机题库替换锚定追问：它照常有效，本文预测判别力大幅下降
9	预测便宜、判断值钱 (Agrawal et al., 2018)	工具压低预测成本，判断成为互补品	它讲 要素价格 ，本文讲 要素归属的可核性	判断由工具生成而本人能推下去：本文判合格，它无从判定
10	结构化行为面试（实务）	过去行为预测未来行为	它问 叙述 ，而叙述是最易被代做的一环	回溯编码区分"要求叙述"与"要求在本人材料上推进"：本文预测只有后者与后续胜任评级相关

十家共同踩着的第二层前提。逐一问"它把什么当作给定"，九行归于同一句：

一个人的能力，原则上可以在他不在场的时候，由他留下的东西来代表。

承认这句不成立，上述各家都须重写其测量端。需要立刻补充一句：**这句前提在过去是真的——它不是错误，是过期的真理（论证见 5.3 节）。**

2.3 最近的一位：现现实务中的交互式口头评估

有一位比表 1 中任何一位都近，且是在本文初稿完成之后的敌意复查中才被找到的。此处如实记录这一遗漏，它是本文写作过程中的一处真实失误。

近两年，英语高等教育中迅速铺开一种称为**交互式口头评估**（interactive oral assessment, IOA）的做法。一所大学的两位教师已连续运行两年：学生就**自己已提交的作业**与教师进行结构化对话，以追问探查其推理而非核对记忆；因问题因人而异、实时展开，故难以外包——机器可产出文本，却撑不住一场关于本人作品的追问（Richter & Dodd, 2025）。另有面向本科生物科学课程的实证研究，报告该形式对包括生成式工具误用在内的学术不端具有抗性（Assessment & Evaluation in Higher Education, 2025）。

必须承认：这套做法与本文第 3.4 节的四条约束几乎逐条同形。若本文的落点停在"应当用这种追问来评价人"，则本文只是这套现现实务的事后理论包装。

分界只有一条，即本文 3.3 节的落点：

那套做法是在验证一个事实；本文说的是构成一份凭据。

表 2 将其展开为四处可判的差别。

表 2 验证与构成：本文与交互式口头评估的四处差别

维度	交互式口头评估	本文	判定形式
待测量	预设存在（这份作业是不是你写的），追问是取样	主张不存在，追问是 生成	若能在人机连续产物上做出可复现、判者间一致的贡献分解，本文第一层即伪（见 11 节 F9）
	学术诚信，默认代做是违规	明写"没有人在造假"，指向结构性解绑	

动因 与人 群			取一批工具使用完全合规、披露完整者：诚信 路径判其无问题，本文预测其中仍有可观比例 落入靶格
作用 域	课程之内；读数进入课程 成绩	课程之外；缺口正在 无课可上的在职者身 上	查其读数是否曾进入可携带凭据（第6节审计 给出否定答案）
时间 结构	产出 签发型凭据 ：通过即 登记，此后不再到期	要求 持有者时间 ：效 力挂在最近一次续签 上	两年前表现优异、今日答不出：它不作处理， 本文判署名已失效

这一发现同时为本文提供了一处正面证据：该做法出现在教育系统而非用人系统，恰与5.4节的三条稳定条件一致——课程内候选人不过剩、判者本来在场、时滞短。**它出现在缺口最浅处，而在缺口最深处（大规模匿名的用人筛选）至今未出现。**

2.4 尚未被占据的位置

综合 2.1-2.3，本文所处的位置可精确表述为三点：（i）不是更好的测评方法，而是关于竞争力**是什么**的主张；（ii）不是诚信问题，而是诚实情形下同样成立的结构性解绑；（iii）不是缺一个字段，而是记录体系在时间结构上的**系统性偏侧**。

3 理论框架

3.1 三个来源领域

为避免与就业研究内部共享过多前提，本文取三个与劳动力市场无关的领域。

3.1.1 把优劣合成一个可携带指数。动植物育种在 20 世纪三四十年代解决过结构相同的难题：个体身上有多个可测性状，价值不一且彼此相关，如何挑出最优者？答案是**选择指数**：按经济权重把各性状合成一个数，使其与总遗传价值的相关最大化，再按此数排序（Smith, 1936; Hazel, 1943）。两条内部主张需记住：其一，**补偿原则**——该法明确允许某项的多余优点抵消另一项的轻微缺陷，“优”因而被定义为可加总、可折算的量，此法在与逐项设卡淘汰的竞争中以效率胜出（Hazel & Lush, 1942）；其二，**该总分挂在个体身上**，因为它要预测“用此个体作亲本，下代能改进多少”，故必须随个体走。其立场：**看显露出来的加权总分就够了**。

3.1.2 把同一顶点按到达方式重新计价。登山界很早发现“登顶了”这句话不含信息量——用直升机上也是上去。于是攀登被划为若干种游戏，每种由一串“不许”定义：不许固定绳、不许膨胀钉、不许补氧、不许提前铺路。两个特征反常识：规则**负面表述**而目的正面——保住根本挑战不被技术侵蚀，技术每让一事变易，规则即多加一条禁令（Tejada-Flores, 1967）；判定权在共同体手中且**可事后追认撤销**——一次极著名的首登声称因拿不出足以服人的路径物证，在数十年后被整个共同体划掉（UIAA, 2002）。其立场：**看它经什么路径做成就够了；显露本身不算数**。

3.1.3 把个体差异按定义清零。群落生态学中的中性理论假设同一营养层内所有物种的所有个体在出生、死亡、扩散乃至成种的人均速率上完全等价；物种在形态、生理上可以千差万别，只要人均生命速率不差，模型即成立（Hubbell, 2001, 2005）。其立场：**看那片条件场就够了；个体高下不必进入模型**。

三者不可通过“判谁对”化解。同一次交付，第一家计出总分，第二家可能计零（若总分靠被禁手段取得），第三家称此分差本不该进模型。三者不在同一语域作答，而共用的标尺并不存在。唯一出路是下探：**它们共同站在什么上面，才使这场争论有意义？**

3.2 共同前提及其反证

共同前提：

一个个体的"优"，是可以脱离具体显露场合、随身携带、跨场结转的东西。

育种将其写成与环境无关的总遗传价值；登山把风格声誉追认给攀登者本人，此后随身；中性理论取消个体间差别，却仍把"等价"按人均安在个体头上。三者争的是这份"优"记在谁名下、算多少，而"它是可有归属、可随身走的东西"被共同视为不必论证。

反证来自最依赖它的那一家自己，且在该领域内部属常识，只是从未被用来回答"优是否随身"。表 3 列出三件材料。

表 3 推翻共同前提的三件材料（均来自第一来源领域自身）

类型	内容	出处
形式化	同一性状在两个环境中，形式上应被当作两个遗传相关的性状，而非同一性状的两次测量；仅当相关系数等于 1 时跨环境结转才成立	Falconer (1952)
实测值	国际种畜评估常年公布跨国遗传相关矩阵，产奶等核心性状的跨国相关显著小于 1；其直接后果不是分数缩水而是 排名重排	Interbull 例行评估；参见 Schaeffer (1994)
制度事实	因重排普遍存在，国际评估于 1990 年代中期放弃全球统一榜，改按各国尺度分别出榜—— 在全球最大的遗传评估体系的正式产出中，"世界最优个体"这一栏不存在	Schaeffer (1994) 及此后的多国评估实践

结论：

"优"不是随身之物，而是配对读数。离开具体条件场，它不缩水，它换算；换算带折损，折损率为 1 减去该相关系数。

方法学声明（双重身份）：本文的分析工具取自第一来源领域，而被推翻的共同前提亦由该领域自身材料推翻。此处有一条硬约束：**用于推翻的只能是它测出来的东西，不能是它主张的东西**。表 3 三件均为观察结果与制度事实；其主张（"优可合成为可携带指数"）正是被推翻的对象，不得同时充当推翻工具。

3.3 从"优不随身"到"可续签的署名"

命题的推导分六步，每步都短。

第一步：若"优"不随身，则随身的只剩一件事——**此人走过某段路**，而路会改写走它的人。这是历史事实，不是存量。

第二步：历史事实自身不显露，历来靠代理显露（文凭、作品集、绩效、推荐信）。

第三步：这些代理共享同一挂点（1.2 节）。

第四步：挂点被一次性剪断，**全部旧代理同批失效**。

第五步：还能显露"路改写过此人"的，只剩一种形式：把他放回自己那件产物的延长线上，当场索取一步未经预备的动作。

第六步（本文的落点）：上述五步只说到“该看哪里”，未说清被看的那个东西是什么。通常的答案是“谁做的变得难以核实了”，据此加强核实——查重、检测、留痕。该答案预设“谁做的”仍是一个已存在、只是暂时看不清的历史事实。

本文主张该预设不成立。一份由人给出方向、由工具生成主体、由人做最后取舍的产物，“它是谁做的”**不存在一个等待被查明的答案**：这里没有隐瞒者，也没有被隐藏的真相；参与是连续的、交织的、逐句不同的。要求指认唯一作者，近似于要求指认一场合唱中“这个和弦是谁唱的”。

由此得到**命题 1（承重命题）**：

人工智能时代大学毕业生的核心竞争力，不是一份能力存量，也不是一次被验证过的作者身份，而是「可续签的署名」——署名是一份没有到期日、却已经会到期的凭据，其效力只能在自己产物的延长线上被当场重做。

推论 1.1：脱稿步不是一种考核手段，而是这份凭据被续签的那个动作本身。这是本文与一切“更好的测评方法”的分水岭：测评方法预设一个待测量，追问只是取样；此处没有待测量——**追问不是在读一个数，它是在生成那个数。**

推论 1.2：读数 q 测的不是能力，而是署名的续签成功率。因此 4.1 节“分母只能是已发起的追问次数”不再仅是技术规定，而是该凭据定义的一部分：一个从未被追问过的人不是 q 为零，而是**其署名从未被要求续签。**

推论 1.3：辨别装置必然是二维的。若署名是事实，一根轴即足够（做没做）；正因署名是凭据，才必然有“凭据在手”与“凭据现在还灵不灵”两件独立的事——档案是前者， q 是后者（第 5 节）。

3.4 定义 1：脱稿步的四条约束

定义 1 脱稿步（off-script step）指同时满足以下四条的一次动作：

（一）锚定本人产物。不是通用题，而是他自己交出的这一件。通用题可刷，锚定追问不可刷——题目的一半由他自己写就，而出题人事先并不知道他会交什么。此条同时保证不可预备性与可判性（判者手中有其产物作为对照）。

（二）在延长线上。不是复述已有内容（那是记忆），而是往前推一步。三种合格形态：换口径重算；往下推一步（若此结论成立，接下来必然出现什么）；指出其在何种条件下会塌。

（三）无脚本。不预告、不给题库。**此条最贵也最脆：**它带来本文的全部判别力，也带来本文最严重的反噬（8.3 节）。

（四）可核。走出的那一步须能被当场判对错或判合理，非印象分。

边界厘清（三种像它而不是它的东西）：压力面试测情绪耐受；刁难细节考记忆——**判法：该问题的答案在提问之前存在吗？**存在者为记忆题，须在应答那一刻被做出者才是脱稿步；“讲讲你的项目”是复述，而复述恰是最易被代做的一环。

3.5 定义 2：签发方时间与持有者时间

命题 1 把竞争力改写为“会到期的凭据”，随之带来一个可查的问题：现行记录体系如何处理“到期”？第 6 节给出审计结果，此处先给出所需的两个概念。

定义 2a 签发方时间：凭据的有效期、撤销与更新，由签发方的日历与行政判断决定。

定义 2b 持有者时间：凭据是否仍然有效，取决于持有者此刻能否将其续上。

命题 2：现行学习与用人记录体系**只有签发方时间，没有持有者时间。**

命题 2 强于“缺一个字段”：缺席可通过增设一栏补上；**偏侧不能**——在签发方一侧增设任意多栏，都不会长出持有者一侧的时间。

3.6 三个来源领域在命题 2 上的分工

三者在此各贡献一块，且不可互换：

- **攀登伦理贡献**“凭据可被事后追认撤销”这一观念——它是本文所知唯一把**可撤销性**写入价值判定的人类共同体；那次数十年后被划掉的首登，是持有者时间的一个原型（不是规则变了，而是**拿不出走过的证据**）。
- **育种的出榜制度贡献**“凭据必须标注作用域”——放弃全球统一榜，等于在凭据上写明“本证仅在此场有效”。
- **中性理论贡献**“存留不需要凭据，承认才需要”——一片条件场可让人长期存活而从不签发任何东西给他。此条同时构成对本文的第一处削弱（第 9 节）。

删去任一家，命题 2 塌一角。

4 测量：脱稿成立率 q

一个不能被清点的判断只是感叹。本节把命题 1 做成可操作的读数。

4.1 定义与计算式

读数名：脱稿成立率 **符号：**q

定义：在一段时期内，针对某人名下产物发起的全部延长线追问中，其应答被判为成立的比例：

$$q = (\text{判为成立的次数} + 0.5 \times \text{判为部分成立的次数}) \div \text{已发起的追问总次数}$$

分母纪律：分母必须是**已发起的追问次数**，不得包含“应当发起而未发起”的次数。后者不可清点；一个从未被追问过的人，其 q **无定义**而非为零（推论 1.2）。空位按定义数不出来，一旦写入分母，该读数即废。

4.2 编码规程（七条）

规程 1（粒度） 一次追问 = 一个锚定问题 + 一次完整应答；中间的追加澄清计为同一次。**边界例：**问“这个跳点是什么原因”，答“口径问题”，再追问“换成下单口径会怎样”——计**一次**（第二问是第一问的完成）；若转向同一份报告的另一张图，计**第二次**。

规程 2（锚定判定） 追问须指向被评者名下具体产物的具体部位。通用情境题、行业常识题、假设性案例题不计入分母。**操作判法：**把该问题原样交给另一位候选人，若他也能答，则不是锚定追问。

规程 3（在场与查阅） 应答须在场完成。**允许查阅**本人原始材料（数据、代码、笔记），不允许离场准备、求助他人或工具。允许查阅的理由：走过路的人知道去哪里查；此条与规程 5 配合，构成全套规程中区分力最强的一对。

规程 4（判定与不一致率） 由两名判者独立给出成立／部分成立／不成立；不一致时按偏严一侧计，并**单独记录不一致率**。不一致率高于 0.30 的场次整场作废——该情形指向追

问本身出得不清楚，而非被评者的问题。（该阈值为设计值，尚未实测；见 6.6 节与第 10 节。）

规程 5（部分成立） “我不知道，但我知道该去看哪里、看什么、大概会看到什么”计部分成立（0.5）。**此条是全套规程中最要紧的一条：**它把“记住了”与“走过了”分开——走过那条路的人即使忘了细节，也知道路的形状；未走过的人连该往哪儿找都说不出。

规程 6（零分与不适用） 拒答、顾左右而言他、把问题重新解释成另一个自己能答的问题，计零。以下三种记**不适用**，不得记零：团队产物中无法界定本人承担部分；涉及保密无法当场展开；被评者当日健康或语言条件明显受限。

规程 7（口径唯一） 本文正文、检验条款（第 11 节）与赌注三处引用的 q，使用同一定义、同一套规程、同一粒度。一个读数若在三处用三种粒度，写死日期的赌注即无法裁决，而不可裁决的赌注等于没有赌。

记录表见附录 A；编码示例见附录 B。

4.3 q 的使用限制

（一）不得跨机构直接比大小。依据来自第一来源领域自身的教训：可遗传性随环境而变，同一量在不同场中不是同一量。q 是**场内读数**——同一机构、同一套判者、同一时段内可比；跨机构比较须先报告场差，或不比。

（二）不得用于排序发薪、一票否决或解雇的单一依据（伦理三条见 8.5 节）。

（三）不适用于尚无“自己名下产物”的人群（如在校低年级学生），锚定追问的前提是有东西可锚。

5 辨别装置与机制

5.1 二维四格

- **横轴：**名下产物档案的厚薄（学历、作品、绩效的存量与光泽）
- **纵轴：**脱稿成立率 q 的高低（署名的续签成功率）

表 4 辨别装置：档案厚薄 × 脱稿成立率

	q 高	q 低
档案厚	① 实走者	② 代走光泽格（靶格）
档案薄	③ 被低估者	④ 常规格

① **实走者：**产物与路径仍绑在一起。**要点：**在代理制年代，此格是**默认**——凡档案厚者几乎必在此格；今天它不再是默认，而是一个**需要被单独确认的结果**。本文的实践含义可压缩为这一句：从前不必确认的事，现在必须确认。

② **代走光泽格：**档案厚而延长线走不出来。按推论 1.3，它不是“造假被抓”的那一格，而是**持有一份已经过期、却无人有权宣告其过期的凭据的那一格**。

③ **被低估者：**路走过了，但未变成可展示的产物——项目烂尾、单位不给署名、缺乏包装资源、岗位不产出成果。**本文对此格作一条明确预测：它是价值重估幅度最大的一格。**理由是对称的：当产物不再能证明能力，产物的缺席也就不再能证伪能力。

④ **常规格：**路没走，产物也没有。判定它不需要本文。

为何靶格是②而非④：一个需要用框架才能看见的靶格才是真靶格。④一眼可见，任何人都需要这套装置。

5.2 靶格的三个签名

签名一：短期指标表现为改善。不是“看起来还行”，而是各项指标真的更好：产出更多、交付更快、格式更规范、文档更完整。管理者看到的是一次成功的效率提升，而该判断在其自身的时间尺度上是正确的。

签名二：失效不产生任何信号。现行流程中没有任何环节负责发起延长线追问：面试问通用题，考核看交付物，评审看成果，晋升看履历。缺口不是“被发现得晚”，而是**没有任何装置会发现它**。后果的形态是：落入②格者可能三年内一路顺利，直到遇到一件必须自己往前推一步的事——一次没有先例的故障、一个客户临时改了口径、一份没有模板可循的判断。届时缺口第一次显影，而此人可能已带团队。**信号来得太晚，且以最贵的形式到达。**

签名三：自我加固。越代走，档案越厚；档案越厚，越无人不好意思追问——资历本身屏蔽追问。问一位新人“这个跳点怎么回事”是正常面试，问一位工作八年者同样的问题会被读作冒犯。**缺口越大，被光泽盖得越严。**

三者合起来解释一个反直觉现象：**该问题会随评价体系越正规化而越严重**——越正规的体系越依赖可归档、可比对、可审计的材料，而这些材料恰恰最易被代做。

5.3 机制：以产物代理路径，在其自身范围内是正确的

指出缺口只完成三分之一。缺口不会自己出现，**它是被某一步操作压掉的**；要证明这是机制而非事故，必须找到那一步并证明**那一步是正确的**。

那一步是：**以产物代理路径**。其正确性有三层，一层比一层硬：

第一层：它便宜一个数量级。

第二层：在产物只能由本人做出的年代，它是**无偏估计**——噪声有，系统偏差无。用一个无偏且便宜的估计替代昂贵的直接观测，是任何工程师都会做的选择。

第三层（最硬）：在那个年代，**直接观测路径反而更差**。路径观测需要判者在场、需要主观判断、需要为每人单独安排，因而引入判者之间的分歧与机会的不平等——谁被资深者带过，谁就被看见。以产物代理路径不仅便宜，而且**更公平、更可审计、更能抵抗人情**。20世纪的考试与文凭制度正是靠这一点战胜荐举制。

机制/事故判别句：把这一步改对，缺口会消失吗？改对它意味着回到“产物是路径无偏估计”的世界，而该世界的成立条件（做出它只能由本人走）已经过去。**改不回去**。因此这不是失误、疏忽或资源不足，而是一个**在自身条件下正确的操作，在条件改变后留下的缺口**。

推论（处方形态）：不应责怪任何环节的执行者。招聘官没有偷懒，教务处没有失职，人力资源部门没有失察——他们执行的是一个曾经正确的设计。**要改的是设计，不是执行**。

5.4 稳定条件与边界条件

缺口能长期存在而不被察觉，需要三条同时成立：

条件 1（供给过剩）：合格候选人多于岗位时，筛选方无动机做更贵的观测。追问是只有在“选不出来”时才值得付的成本。

条件 2（成本无科目）：一次锚定追问需判者事先阅读被评者产物，通常 15–30 分钟，两名判者则加倍；这笔时间在多数机构的招聘预算中**没有对应科目**，因而永远排在最后。

条件 3（反馈时滞）：错判的后果通常在 2-4 年后显露，届时决策者可能已换人，反馈回不到评价环节。**没有反馈的系统不会自我纠正。**

边界条件（三条全部不成立之处）：熟人劳动市场——师徒制、家族企业、小镇诊所、乡村学校。那里候选人不过剩、判者天天在场、时滞极短。**因此那里不存在本文所说的缺口：路径不需要被记录，因为它从未离开视野。**

由此得到一条尖锐的推论：**本文所描述的缺口是大规模匿名筛选的副产品，而非人工智能的副产品。**人工智能只是把它从潜在变为现实——它一直在那里，只是从前有一根绳子替我们兜着。

6 研究设计与三项字段审计

6.1 总体设计

命题 2 是一个关于现存文档的经验主张，可用最低成本检验：**去查账本里有没有这个字段、失效由哪一侧触发。**本文执行三项审计，均在读取数据之前写死三件：判负条款、不利结果处置、停止规则。

表 5 三项审计的设计与结果概览

	审计一	审计二	审计三
对象	学习记录互操作标准 A（1EdTech CLR 2.0, 2024-04-02 候选终版）	同上（同一份对象）	学习记录互操作标准 B（欧洲学习模型 ELM，本体 v3.2.0）
来源	美国主导的非营利标准组织	同上	欧盟委员会（EMPL 总司）
单位	共享凭证数据模型 32 个类	同上	本体 52 个类及其对象/数据属性
检验问题	有无记录该动作的带类型字段	失效触发条件在哪一侧	两个问题同时检验（独立复算）
结果	0/32 命中判负条款； 强主张按自设条款判负	0/32；三类失效机制全在签发方一侧	0/52； 独立复算成功 ，两个结论均复现
独立性	—	与审计一共用对象， 非独立样本	与前两项来源、机构、法域、技术谱系均不同

6.2 审计一：是否存在带类型的字段

事先写死：（判负条款）若在标准中找到**任一具名字段**，其定义要求记录"学习者在其自身既有产物延长线上的一次即时、非预设的应答"，则"账本无此字段"判负。（不利结果处置）无论方向如何原样写入，不改主张再跑第二份。（停止规则）只查该标准正式发布版本的共享凭证数据模型这一份清单，逐类查完即停。

结果：0/32 命中。三个最接近者均不是：**证据类**定义为"支持某项主张的信息，例如指向学习者所产出物件的链接"——记的是**物件**；**成就类**挂"达成标准"，属**预先写定**的条件，与定义 1 第三条直接冲突；**结果类**挂评分量表与等级描述，同样预设。

两处不利结果，原样写入。其一，该标准设有背书机制：当个人对自身知识或技能做出自我断言时，认可机构可出具背书表示同意，背书凭证带有与成就凭证大体相同的元数据并允许附加文字说明——**这是一个可以把脱稿步塞进去的自由文本插槽。**其二，该标准明写数据模型可扩展。

因此"账本里没有这个字段"这一强主张，按本文自设条款判负。修订后仍成立的是：

账本中没有为它设置的、带类型、可比较、可加权、可被查询的字段。它只能以非结构化叙述附着于其他字段，因而进不了任何排序、筛选与绩效计算。

此句弱于原句而更准，并把本文推到更硬的位置：**缺的不是存储空间，是可比较性。**一件事只要不能被比较，在任何选拔制度中即等于不存在——哪怕它被完整记录。

6.3 审计二：失效的触发条件在哪一侧

事先写死：（判负条款）若存在任一字段，其失效或撤销的触发条件绑定于持有者当下的一次表现，则"失效只由签发方一侧触发"判负。（处置与停止规则同上。）

先报告不利于本文初稿的发现：该标准有失效机制，且不止一处。第一类，**有效期**——断言元数据带签发日与到期日；第二类，**凭据状态**——数据模型设有专门的类表达凭据当前状态，撤销通过它表达，用例中亦明确写到证照被吊销或失效的情形；第三类，**刷新服务**——持有者钱包可据此回到签发方拉取最新版本。故"记录体系里根本没有到期这回事"这句话是错的，此处划掉。

但三类并排即见其共同点：触发条件全部落在同一侧。有效期由签发时预设的日历决定，与持有者此后所为无关；撤销由签发方的行政判断触发（发错了、资格取消了、政策变了）；刷新是回到签发方去问最新状态。

逐类核对：0/32 命中判负条款。

6.4 审计三：在一份来源独立的标准上复算

动机：审计一与审计二共用同一对象，**不是两个独立样本**。为消除这一缺陷，本文在一份来源、机构、法域与技术谱系均不同的标准上重跑同样两个问题。

对象：欧洲学习模型（European Learning Model, ELM）本体 v3.2.0，由欧盟委员会就业、社会事务与包容总司下属单位维护，为欧洲数字学习凭据（EDC）、资历数据集登记（QDR）等应用剖面提供统一词汇；官方说明称其在各应用剖面上共有五百余项属性。本次审计单位为本体列出的 **52 个类**及其对象属性与数据属性清单。

事先写死：判负条款、不利结果处置与停止规则与审计一、二逐字相同，仅替换对象；停止规则限定为"只查本体文档所列的类与属性清单，查完即停，不追加应用剖面的 SHACL 约束文件"。

结果 A（对应审计一的问题）：0/52 命中。最接近的三处均不是：

- **证据类（Evidence）** 定义为"关于导致该可验证凭据签发之过程的信息"。**这比标准 A 的对应类更接近本文所需**——它明确指向"过程"而非物件。但它指向的是**签发过程**（该凭据是如何被授予的），而非持有者在此后某一时刻于自身产物延长线上的应答。**这是三项审计中距离判负条款最近的一次，仍未命中。**
- **学习评估类（LearningAssessment）** 挂有"评估方""评估状态""成绩""成绩分布""简化评级"等属性，其中**"身份核验"（ID verification）**属性最值得注意：它记录评估过程中学习者身份是否经过核验。**这是一处极有价值的近失**——标准已为"是不是本人"设了字段，却未为"本人此刻能否续上"设字段。按 2.3 节的区分，前者属**验证**，后者属**构成**。
- **备注类（Note）** 为自由文本，与标准 A 的背书插槽同属可容纳而不可比较的一类。

结果 B（对应审计二的问题）：0/52 命中判负条款，且失效机制族更完整。本体中与失效相关的至少有五处：**到期日**（expiryDate 数据属性）；**凭据状态类**（定义为"关于某可

验证凭据当前状态的发现信息")；**权利状态** (entitlementStatus，挂在学习权利规范上)；**审查日期** (reviewDate，挂在认证/质量保证上)；**核检检查类及其状态属性**。逐一考察其触发方：到期日由签发时写定；凭据状态由签发方维护；权利状态与审查日期由主管或质量保证机构掌握；核检检查是**验证方**对凭据完整性的检查，检查对象是凭据本身而非持有者的表现。

结论：两份来源独立的标准，在两个问题上给出一致结果。命题 2 由此获得一次独立复算：

现行学习记录体系只有签发方时间，没有持有者时间。

代理变量与覆盖率：两份标准均属学历与学习凭据一侧；**企业人力资源数据交换标准与各国学历核检系统尚未审计**，留作复算（见第 11 节 F10）。

6.5 三项审计共同的方法局限

本文读取的是标准的类定义、术语表与属性清单，**不是每个类的全部属性的逐条穷举**，也未读取各应用剖面的约束文件。若某类的某个属性带有本文所需语义而未在上述文档中体现，本次审计会漏掉它。**这是一处真实的分辨率不足，不作辩解**。其方向是保守的还是激进的，本文无法先验判断：漏掉一个字段会使结论偏向本文，因此第 11 节 F10 将复算列为优先待办。

6.6 研究四（预注册的待做检验）：判者一致性

规程 4 中 0.30 的不一致率阈值为**设计值而非实测值**。本文在此给出一份可直接执行的预注册方案，并明确声明**该研究尚未实施**，其结果不进入本文任何结论。

设计：在单一机构内，招募两名经培训的判者，对同一批被评者的同一批锚定追问独立评级（成立／部分成立／不成立）。追问由第三方按定义 1 生成，判者互不通气。

主要终点：两判者之间的加权 kappa 一致性系数。

样本量：以 Donner & Eliasziw (1992) 的方法，取零假设 $\kappa_0=0.40$ 、备择 $\kappa_1=0.70$ 、 $\alpha=0.05$ （双侧）、检验功效 0.80、三分类近似等边际分布，所需追问次数约为 110 次；按每位被评者 3 次追问计，约需 37 人。**本文取 N=120 次追问（40 人）以留出不适用与作废场次的余量。**

判负条款（事先写死）：若加权 kappa 的 95% 置信区间下界低于 0.40，则规程 4 不可用，脱稿成立率 q 在当前编码规程下不具备可复现性，第 4 节须整体重写。

为何这是本文最优先的一项：q 的全部合法性依赖于两名判者能否对"成立"给出一致判断。若不能，本文的贡献二与贡献三同时塌陷，而贡献一与贡献四不受影响（分层引用见 13 节）。

7 五类制度场景

第 3.5 节的两类时间是否有现实对应？本节以一句不含任何程度词的判据询问五类制度：

在这家机构的录用与考核流程中，"交付合格之后、锚定本人产物的当场追问或变体重做"被记录在哪个字段、占多少权重？

五个答案互不相同，其中两个（场景四、场景五）**推翻了本文的先验预期。**

表 6 五类制度场景的判据答案

场景	追问是否发生	是否进入字段	是否可携带	时间类型
一 · 高校学位答辩与课内口头评估	是	部分（见下）	否	签发方
二 · 公务员考录	否（非锚定）	是（面试分）	否	签发方
三 · 软件业代码审查	是	否	否	签发方
四 · 民航飞行员定期检查	是	是	是（绑定执照）	持有者
五 · 学术同行评审	是	是（对稿件）	不适用	对象不是人

场景一 · 高校学位答辩与课内口头评估。答案分两半，第二半推翻了本文初稿。前半：学位答辩中追问确实发生（委员会当场提问、锚定本人论文），但**追问表现不进入成绩单的任何字段**——成绩单上只有“通过／不通过”与总评等级，追问内容与应答质量留在答辩记录中，此后不再被任何人读取。后半：2.3 节所述的交互式口头评估有**评分量表并计入课程成绩**。因此“答辩类追问一律不进字段”这句话是错的，此处划掉。修订后的准确表述：**它进入课程成绩，不进入任何可跨机构携带、可查询的凭据**——按定义 2，它产出的仍是签发凭据，一次登记，此后不再到期。

场景二 · 公务员考录。面试环节权重公开（多地占最终成绩的一半），但测评要素为通用素质清单（综合分析、计划组织、人际沟通、应变），**不锚定考生自身的既有产物**。按规程 2，此类一律不计入分母。**权重很高，锚定为零。**

场景三 · 软件业代码审查。这是既有实践中最接近脱稿步者——审查者对着此人提交的这段代码提问，问题往前推（“并发情况下这里会怎样”“这个边界值你试过吗”）。但其记录形式为逐条评论与“通过／需修改”，**没有任何字段记录提交者当场回应的质量。动作齐全，读数被扔掉。**

场景四 · 民航飞行员定期检查。这是五类中唯一把持有者时间写入正式字段者：周期性模拟机复训与检查中，检查员临时植入未预告的异常状态，考察机组处置；结果以“合格／需补训／不合格”记入个人技术档案，并直接绑定执照有效性。**本文原本预期没有任何场景会命中，而它命中了。**它同时给出边界条件：该字段之所以能存在，依赖失效后果极端、监管强制、成本全行业分摊三者同时成立——同时具备这三条的场合，全社会大概不超过几个。

场景五 · 学术同行评审。追问发生（审稿意见锚定本稿），但它测的是稿件而非作者——作者身份常被匿名化，评审结论不写入任何关于作者能力的记录。**追问的对象是产物，不是人。**此条反过来给本文加了一条限制：一个成熟的、以追问为核心的制度，可以完全不产生关于人的读数。

五个答案的分布本身即是一项结果：不进字段（一）／权重高但不锚定（二）／动作有而读数被扔（三）／完整写入字段（四）／对象不是人（五）。只有第四个是完整的，而它是**社会成本最高的那一个。**

8 讨论

8.1 三项发现的合并含义

三项审计给出一条比“缺一个字段”更硬的结论：**记录体系在时间结构上系统性偏侧。**这一结论解释了三件此前解释不了的事。

其一，为什么学历越老越不被追问。在签发方时间里，一份 2018 年的学位与一份 2025 年的学位效力完全相同——两者都未到期、都未被撤销。年份在这套时间里只是属性，不是损耗。资历因此天然屏蔽追问（签名三的机制来源）。

其二，为什么“增设一门课”这类对策注定无效。它增发的是新凭据，而问题不在凭据不够多，在已发出的凭据无法在持有者一侧到期。

其三，为什么场景四是唯一完整的一例。它是本文查到的唯一把执照有效性绑定在周期性、未预告的当场处置表现上的制度，而其存在条件苛刻到几乎不可复制。

8.2 对三方的启示

对毕业生。其一，保留一条自己走过的路，并知道它在哪儿：不必多，一件足矣，包括做错的部分、绕过的弯路、最后没解决的那个问题。**烂尾的项目不是污点，而是延长线最长的那种材料**——因为你知道它卡在哪里，而卡住的地方是代做不出来的。其二，**用工具，但保留判断的落点**：判断的痕迹不在结论里，在**被放弃的那些选项里**。其三，**主动求追问**：缺口在被雇佣之前暴露，代价最小。

对大学。本文的处方不是加课，而是在**已有环节上装一个读数口**。场景一显示，答辩中的追问已经发生，只是读数被扔掉；把它捡起来的成本接近于零。课程层面最便宜的改动，是在既有作业与实验中加一个延长线环节：交完之后随机抽取一部分学生，就其自己交的这份问一个往前推的问题——**抽样即可，不必普测**（理由见 8.3 节）。

对雇主。场景三显示，代码审查已在做锚定追问，只是不记读数；最便宜的改动是把已有动作的结果记下来，而非新增流程。招聘中把“讲讲你的项目”换成“就你交的这份材料，往前推一步”，同样时间，判别力完全不同。晋升中须格外小心：**q 越往高层用越危险**——签名三说明高层最不习惯被追问，而对高层追问最易被读作羞辱；建议仅在新岗位交接时刻由交接双方共同参与，而非作为定期考核。

8.3 反噬预言：最省事的实现方式会毁掉它

把延长线追问标准化成一套题库。

这几乎必然发生，且不因有人心怀恶意，而因机构需要**可比、可培训、可审计**的东西，而无脚本恰恰**不可比、不可培训、不可审计**。一个不能培训的环节，在任何有培训部门的机构里都活不过两年。题库一出，次日即有人开始刷；刷完之后，脱稿步变回脚本步，判据自毁。

这不是副作用，而是本文的核心性质：它的判别力全部来自“不可预备”，而任何制度化努力都在制造可预备性。

更糟的一层：该读数一旦进入考核，最省事的达标办法通常是在文书上做手脚——把常规问答记成“延长线追问”，把 q 做成人人 0.9 的数字。届时它既不再判别任何东西，又占着一个正式字段。

本文看不到出口。能想到的最好办法只是延缓：**抽样而非普测**（普测必然催生题库）、**判者具名并轮换、追问文本随记录留存且可被申诉、明令禁止公布题目样例**。因此本文反对将其写入任何普遍适用的评分表，亦反对作为全行业标准推广。**它应以抽查形式存在，而非以指标形式存在**。一个不能被普遍化的判据在制度上是笨拙的；而它笨拙的原因恰是它有效的原因。

8.4 与同一批两项研究的关系

本文与另外两项处理同一问题的研究互为最近邻，须给出分界。

与《现配优势》（该文主张：竞争差不是存量，而是每次交手由可传构件加一手不入账的配组当场做出的显露；读数为**留种率**，即把完备交接清单交给同资历者执行后原表现的复现份额）。**分界**：该文锚在**人与场之间**，本文锚在**产物与走出它的路之间**。**判决性对照**：取一批档案很厚、且**从未被追问过**的人——该文仍给出读数（交接清单复现份额），**本文不给读数**（q 在追问未发起时无定义）。反之，若观察到“清单交接后复现率高、而延长线追问答不出”的人群稳定存在，**本文对而该文错**：这说明优势可被交接（非现配），而此人仍不是走出它的那个。**互补**：该文解释“带走一个人不等于带走表现”，本文解释“看一份产物不等于看见做它的人”；两者叠加时后果相乘。

与《零判份》（该文主张：一次合格产出中有一部分判别力为零——不是低质量、不是抄袭，可能恰是最原创最费力的一段；读数为**未见率**，即交付日一个本科水平者用公开资源在 30 分钟内能否产出功能等价物）。**分界**：该文锚在**产出与世界现存样本之间**，本文锚在**产物与路径之间**。二者测同一份产出的两件不同的事：未见率高而并非由署名者走出，或未见率为零而署名者确实一步步走过，皆可能。**判决性对照**：取一人在同组织、同配组、完全不变条件下连续五年的交付——该文预测判别力单调下降，**本文对此情形不作预测**（q 只在追问被发起时有读数）。**这是三项研究间最清晰的一条分界：前两项读产出，本文读一次被发起的追问**。若三者在同一份台账上给出相反评价，取发起过追问的那一份为准——因为另两者都在读一个已不能代表人的东西。

一处坦白：三项研究都在处理“能力的显露出了问题”，因而不是三个独立发现，而是同一件事的三个切面。本文新增的部分只有一件：**把缺口定位在记录端而非能力端**，并给出可清点的追问读数。

8.5 伦理

本读数会被拿去考核人，故三条须写死。**三条缺一，本文的处方不应被采用。**

其一 · 该读数指的是位置，不是人。q 很低者，可能只是从未有人给过他一次可以往前走的追问，也可能其岗位从不产生“自己名下的产物”。**把 q 读成“此人不行为”是误用**——它首先是对机构的诊断。

其二 · 不得为使该数好看而制造追问。尤其不得在有安全后果的岗位上安排突击追问以充实样本。**读数服从工作，不是工作服从读数。**

其三 · 据此做出不利安排的机构，必须公开编码规程并给出复核通道。当事人有权要求换判者重做，有权查看追问原文与两位判者各自的判定，有权对“锚定是否成立”提出异议（按规程 2，一个另一位候选人也能答的问题不算锚定追问）。

附加一条（常识而非伦理）：q 不得用于本科低年级、不得作为实习生转正的一票否决、不得作为解雇的单一依据。抽查读数不承担一票否决的重量。

9 三处反过来削弱本文的约束

一篇文章若三个来源都只在支持它，则这三个来源多半是被挑出来的。以下三处**真的削弱了本文**，每处都削掉了本文原本想说的一句更强的话。

其一，来自中性理论：本文不得声称脱稿步决定谁能就业。存留大量由行业周期、地域、家庭、时机、第一份工作的平台决定，这些与 q 无关。**被削掉的更强主张：**“脱稿步是竞争力本身”；现只能说“脱稿步是承认的凭据”。**适用范围削去一半——**本文管“谁被认作是他自己”，不管“谁能活下来”。实际后果须写明：**q 很高者可能仍找不到工作，q 很低者可能一路顺风。**本文不预测就业结果，只预测“当一次锚定追问被发起时会发生什么”。

其二，来自攀登伦理：判据必须由具名共同体持有，并可事后追认撤销。这直接毁掉本文最自然、最好卖的处方——做一个自动化追问装置给所有人打一个 q。理由是硬的：判据一旦交给不具名、不可抗辩的装置，立刻退化为新题库，而题库可刷；被误判者**没有可申诉的对象**。被削掉的是**价值排序**：本文宁可要慢的、抽样的、可申诉的人工判定，也不要快的、全覆盖的自动判定。**这与效率方向相反**，而效率恰是此类方法最易被采纳的理由。本文在此放弃了它最大的一块市场。

其三，来自育种的实测：可遗传性随环境而变。被削掉的更强主张："q 可以成为一个通用分数"。这是此类读数最有诱惑力的一条路——一个全球可比的数字。本文关掉了它：跨场比较必须先报告场差，或不比。

10 局限

(一) 审计的分辨率不足。见 6.5 节：三项审计读的是类定义、术语表与属性清单，非属性穷举，亦未读应用剖面的约束文件。

(二) 审计的覆盖率不足。两份标准均属学历与学习凭据一侧；企业人力资源数据交换标准与各国学历核验系统尚未审计。

(三) 判者一致性未经实测。规程 4 的阈值为设计值；第 6.6 节的预注册研究尚未实施。**这是本文当前最重的一笔欠账——q 的全部合法性系于此。**

(四) 核心因果主张无个体层面追踪数据。q 与后续胜任之间的剂量—反应关系 (F1) 尚未检验。

(五) 五类场景中有两类来自公开资料的二手描述（场景四、场景五），未做一手核对；若其字段实况与本文所述不符，表 6 的分布结论须改。

(六) 本文无法排除选择效应。q 高者是否本来就更愿意接受追问？观察数据不能分离意愿与能力，需随机指派追问的设计。

11 证伪条件与一条写死日期的赌注

11.1 证伪条件

F1 (主检验 · 机制证伪) 最强的竞争解释是"这不过是面试老手对生手"（表达能力与临场经验）。故：**控制表达能力评分（结构化沟通量表）与既往面试次数后，若 q 与入职 12 个月后的实际胜任评级之间不存在剂量—反应关系，则本判断为伪——届时应归因于一般临场表达能力，而非"路径是否改写过此人"。**样本纪律：同一制度环境内 $N \geq 6$ 个单元（同一公司多部门、同一医院多科室、同一学校多院系），不得以两个国家或两家文化迥异的公司充数。

F2 (顺序而非相关) 在同一机构内，**追问环节的设立应早于用人偏好的转移。**若观察到偏好先变、追问后设，则本文因果方向写反。

F3 (锚定必要性) 以随机题库替换锚定追问，若判别力不显著下降，则定义 1 第一条为伪，本文与挑战应答式认证的分界（表 1 第 8 行）随之失效。

F4 (部分成立档) 若规程 5 的 0.5 档与"完全答不出"在后续胜任评级上无差异，则规程 5 为伪，本文的核心区分（走过 vs 记住）塌掉一半。

F5 (最便宜的一条) 任取一家已做结构化面试的机构，从其**既有面试录音**中回溯编码：每个问题是"要求叙述"还是"要求在候选人自己材料上推进"，再与入职后评级对照。**该数据在多数机构的既有记录中已经存在，一人两周可跑完。**这是本文最应先跑的一条。

F6 (判者一致性) 见 6.6 节判负条款。

F7 (反向条件) 若在熟人劳动市场中亦测到同样的 q —胜任关联, 则 5.4 节的边界写错, 本文的机制解释应让位于更朴素的解释 (追问只是好的面试技巧)。

F8 (杀死条款) 若在三个不同行业各取一批人, 其档案厚度与 q 的相关系数稳定高于 0.7, 则②格在经验上近乎不存在, 靶格选错, 全文作废。

F9 (针对贡献一) 若能在人机连续参与的产物上, 给出一套可复现、判者间一致的贡献分解方法, 把"哪一部分是此人做的"稳定查出, 则命题 1 为伪, 本文应退回"更好的验证方法"这一族。**这是本文最愿意被推翻的一条:** 若被推翻, 问题将有一个远比追问便宜的解法。

F10 (针对命题 2) 在企业人力资源数据交换标准、各国学历核验系统中任取一份, 若查出任一字段其失效触发条件绑定于持有者当下的一次表现, 则命题 2 收窄为"该标准之外成立"。场景四已说明此类字段并非不可能, 只是极罕见。

若被证伪, 我们学到什么: 若 q 与胜任无关, 则说明产物与路径的解绑并未真的发生, 或市场已找到别的代理。**那个代理是什么, 会是比本文重要得多的一个发现。**

11.2 赌注

赌注: 到 2029 年 12 月 31 日之前, 至少有三家员工数超过一万人的机构, 在其公开的招聘或晋升规程中写入一个**锚定候选人本人既有产物的当场追问环节**, 并说明其在最终决定中的权重。

不算命中者 (须写死, 否则赌注不可裁决): 泛称的"结构化面试""行为面试""案例面试"不算——必须锚定候选人自己提交的材料; 仅存在于内部文件、未公开规程者不算; 仅要求"讲述你的项目"者不算——讲述是复述, 不是延长线; 机构自行声明"我们很重视这一点"而无可执行条款者不算; 现场编程测试、笔试、案例分析不算——它们锚定的是题目, 不是候选人自己的既有产物。

作者押其命中。**若到期未命中, 本文的判断即便在道理上成立, 也说明它挑错了显露形式——**届时应回去问: 制度实际找到的替代显露是什么?

12 自反性说明

本文是一件产物。

按其自身主张, 产物不显露路径, **故本文没有资格用自身证明自己**。它现在躺在读者面前: 结构完整、有表、有数、有事先登记的检验、有自我推翻的记录。这些恰是靶格的签名——短期指标全面改善, 形式审查全部通过。

诚实的做法只有一个: **把它交给延长线追问**。此处先给出三问, 供人验:

问一: 把 6.4 节的审计对象换成企业人力资源数据交换标准, **结论要改哪一句?** (提示: 该族标准以岗位与胜任力为中心, 本文预期其字段更接近"能做什么"而非"走过什么", 故修订版主张不变; 但若其带有"评估记录"类的结构化字段, 6.2 节的修订版须再收窄一次。)

问二: 把规程 5 ("不知道但知道去哪儿看"计 0.5) 改为计零, **表 4 的哪一格会松动?** (提示: ③被低估者一格会大幅缩小, 因为薄档案者最常见的合格形态正是这一档。)

问三: 5.4 节称缺口是"大规模匿名筛选的副产品", **那么在完全匿名的开源协作中, 为什么反而看得见人?** (提示: 那里的产物自带过程——提交历史、讨论记录、被拒绝的方

案。这是一个对本文有利的观察，但它同时暗示了一条本文没有走的路：让产物重新带上过程，可能是比追问更根本的解法。此条记为欠账，见 13.3 节。)

答不出来，本文即落在自己的靶格里，且本文自己不会发出任何信号。

另有一处更硬的：本文主张判据必须由具名共同体持有并可追认撤销，故本文自己亦必须可被追认撤销。若 11.2 节的赌注被打脸，本文不作修补，作废重写。

13 结论

13.1 主要结论

我们用一个世纪把"看人"这件昂贵的事替换成"看东西"这件便宜的事。该替换在当时是正确的：它使大规模、匿名、抗人情的选拔成为可能，其社会收益怎么估计都不为过。

现在东西与人之间的那根线松了。松了之后，全部旧办法同时失灵——不是某一种失灵，而是同一根线上挂着的所有办法一起失灵。

剩下的只有一个很笨的办法：把人家叫到他自己做的东西跟前，问他一句他没准备过的话。它慢、贵、不可规模化、不能自动打分、且随时可能因被制度化而自毁。但它是仅剩的那一个，而它笨拙的理由与它有效的理由是同一个：它要求那个人当场在场。

比办法更要紧的是：我们一直以为问题在于"看不清是谁做的"，于是把力气花在看得更清上。这条路走不通，不是工具不够好，而是那个要被看清的东西已经不在。一旦承认这一点，署名就从"曾经发生过、可被核实"的事，变成"必须不断被重做"的事。它有了时效，而我们的账本仍是永久台账：签发日写得清清楚楚，撤销权攥在签发方手里，唯独没有一栏留给持有者——没有一栏，用来记这个人今天还能不能把它续上。

补这一栏，是本文唯一的请求。

13.2 三层引用

本文可分层引用，读者不必全盘接受：

- **第一层（最强、最易倒）**：署名不是可核实的事实而是需要续签的凭据；脱稿步即那次续签。**"不是事实"是最脆的一处**（见 F9）。
- **第二层（分析工具）**：表 4 的四格与"代走光泽格"这一诊断。即使第一层被推翻，该表仍可用于分类。
- **第三层（最稳、最易检验）**：现行记录只有签发方时间，没有持有者时间——失效字段整整一族，触发条件全在签发方一侧，已在两份来源独立的标准上复算。前两层若都被推翻，这一层仍站得住，且它本身已是一件值得动手去补的事。

13.3 四个本文解释不了的洞

1. **q 高者是否本来就更愿意接受追问**？若选择效应主导，则测到的可能是意愿而非能力（见局限六）。
2. **落入②格者会否在岗位上很快把路补上**？若三个月内补齐，则②格只是短暂错配，而非结构性缺口。本文无追踪数据。
3. **机器能否学会脱稿**？若一个模型可实时接管某人的延长线追问并答得成立，则定义 1 第三条失去技术前提，本文全部结论须重写。**作者认为这是最可能先发生的一条。**

4. **让产物重新带上过程，是否比追问更根本？** 开源协作中的提交历史、讨论记录、被拒绝的方案使产物自带路径。若这条路走得通，本文所主张的追问只是一个过渡方案。**这是最大的一笔欠账。**

声明

伦理声明：本研究不涉及人类受试者，三项审计的对象均为公开发布的技术标准文档。第 6.6 节的待做研究若实施，须先获所在机构伦理审查批准，并遵守 8.5 节三条。

利益冲突：作者声明无利益冲突。本文不受任何标准制定组织、评测服务提供商或人力资源技术企业资助。

数据可得性：三项审计的对象均为公开文档，任何人可按 6.2–6.4 节的停止规则原样复算。审计过程未产生除本文表格之外的中间数据集。

工具使用声明：本文的论证结构、承重命题、读数定义与编码规程、三项审计的事先登记条款与判定，由作者负责；文献检索与文字整理由人机协作完成。**三处不利结果（6.2 节的强主张判负、6.3 节的失效机制存在、7 节场景一的自我推翻）未作任何修饰。**按本文自身的主张，此声明不能证明本文的署名成立——它只能在延长线上被重做（第 12 节）。

附录 A 追问记录表（可直接使用）

字段	说明	示例
被评者代号	匿名化	C-021
产物标识	具体到部位	实习报告 v3 · 第 4 节图 2
追问序号	同一场次内连号	2
追问文本	原话记录，不得事后润色	"把统计口径换成下单时间，这个跳点还在不在？在的话第二个结论要改哪一句？"
应答摘要	判者记录，不超过 100 字	指出跳点会消失，并指出第二条结论中"季度末冲量"一句须删
是否查阅原始材料	是/否（允许，仅作参考记录）	是
判者甲	成立/部分/不成立	成立
判者乙	同上	部分
计分	不一致按偏严一侧	0.5
不适用标记	规程 6 三种情形之一	—
备注	场次条件、异常情况	判者乙认为口径定义不清

场次汇总另记三项：**追问总次数、不一致率、不适用次数**。三项缺一，该场次的 q 不得引用。

附录 B 编码示例（四例，含两处边界）

#	追问	应答	判定	依据
1	就本人报告的口径变更提问	当场指出跳点会消失并改正结论	成立 (1.0)	定义 1 四条齐
2	同上			规程 5：知道路的形状

		"我不知道，但这要看订单表的时间戳字段，改了之后季度末那一段会平掉"	部分成立 (0.5)	
3	"请讲讲你这个项目的亮点"	流畅讲述	不计入分母	规程 2：叙述题非锚定追问
4	"你报告第 17 页第三行的数字是多少"	准确复述	不计入分母	定义 1 第二条：答案在提问前已存在，属记忆题

附录 C 三处工具粗糙处的登记

诚实要求把方法的钝处写出来，而非等审读者发现。

其一 · 审计的分辨率。见 6.5 节。**未做逐属性穷举，是分辨率不足，不是覆盖率不足；覆盖率不足另见局限（二）。**

其二 · 判者一致性未实测。规程 4 的 0.30 阈值是设计值。真实判者分歧可能远高于此，届时该读数的可用性须重估。**这是最需要先做的一次预实验（6.6 节）。**

其三 · 五类场景中两类为二手描述。场景四与场景五未作一手核对。

参考文献

1. Agrawal, A., Gans, J., & Goldfarb, A. (2018). Prediction Machines: The Simple Economics of Artificial Intelligence. Harvard Business Review Press.

2. Becker, G. S. (1964). Human Capital: A Theoretical and Empirical Analysis. University of Chicago Press.

3. Donner, A., & Eliasziw, M. (1992). A goodness-of-fit approach to inference procedures for the kappa statistic. Statistics in Medicine, 11(11), 1511–1519.

4. European Commission (2024). European Learning Model (ELM) Ontology, Revision 3.2.0. Directorate-General for Employment, Social Affairs and Inclusion. <http://data.europa.eu/snb/model/elm>

5. Falconer, D. S. (1952). The problem of environment and selection. The American Naturalist, 86(830), 293–298.

6. Falconer, D. S., & Mackay, T. F. C. (1996). Introduction to Quantitative Genetics (4th ed.). Longman.

7. Hatano, G., & Inagaki, K. (1986). Two courses of expertise. In H. Stevenson, H. Azuma, & K. Hakuta (Eds.), Child Development and Education in Japan (pp. 262–272). Freeman.

8. Hazel, L. N. (1943). The genetic basis for constructing selection indexes. Genetics, 28(6), 476–490.

9. Hazel, L. N., & Lush, J. L. (1942). The efficiency of three methods of selection. Journal of Heredity, 33(11), 393–399.

10. Hirsch, F. (1976). Social Limits to Growth. Harvard University Press.

11. Hubbell, S. P. (2001). *The Unified Neutral Theory of Biodiversity and Biogeography*. Princeton University Press.
12. Hubbell, S. P. (2005). Neutral theory in community ecology and the hypothesis of functional equivalence. *Functional Ecology*, 19(1), 166–172.
13. 1EdTech Consortium (2024). *Comprehensive Learner Record Standard, Version 2.0 (Candidate Final Public, 2 April 2024)*.
14. Polanyi, M. (1966). *The Tacit Dimension*. Doubleday.
15. Richter, S., & Dodd, P. (2025, October 2). The case for oral exams in the age of AI. University of Auckland. (另见 *The Conversation*, 2026 年 1 月转载)
16. Schaeffer, L. R. (1994). Multiple-country comparison of dairy sires. *Journal of Dairy Science*, 77(9), 2671–2678.
17. Schmidt, F. L., & Hunter, J. E. (1998). The validity and utility of selection methods in personnel psychology. *Psychological Bulletin*, 124(2), 262–274.
18. Smith, H. F. (1936). A discriminant function for plant selection. *Annals of Eugenics*, 7(3), 240–250.
19. Spence, M. (1973). Job market signaling. *The Quarterly Journal of Economics*, 87(3), 355–374.
20. Tejada-Flores, L. (1967). *Games climbers play*. Ascent (Sierra Club), May 1967.
21. ten Cate, O. (2005). Entrustability of professional activities and competency-based training. *Medical Education*, 39(12), 1176–1177.
22. Thurow, L. C. (1975). *Generating Inequality: Mechanisms of Distribution in the U.S. Economy*. Basic Books.
23. UIAA (2002). *Tyrol Declaration on Best Practice in Mountain Sports*.
24. World Economic Forum (2025). *The Future of Jobs Report 2025*.
25. Zahavi, A. (1975). Mate selection—a selection for a handicap. *Journal of Theoretical Biology*, 53(1), 205–214.
26. *Assessment & Evaluation in Higher Education* (2025). Utilising one-on-one interactive oral assessments as the major final assessment within a bioscience course. DOI: 10.1080/02602938.2025.2502577.

本文为一次性写就的学术论文稿，全文机检 19,337 汉字。作者不为本文自行印发任何评分。

本体层回答了"发生长什么样"，立刻引出一个前四章都不回答的问题：这样的能力靠什么活着？它不是认证过就封存的资产——第二部已经证明它须被续签；那么续签的本钱从哪里来？

《底流》独占这一层，因为它做的事与其余九章方向相反：九章都站在**读数侧**（能力怎么被读出），它换到**供给侧**（能力靠什么被喂养）。它的发现毫不浪漫：喂养发生学能力的，从来不是任何专款——不是培训预算、不是刻意练习的课时、不是传承制度的补贴——而是日常工作里最低价值那部分的联合副产：无辅助、有真实赌注、亲手完成的产出流。正因为它从来没被记过账，它被智能系统按面值接走时，三本账全在改善，断供无声。

本部只有一章，但它在全书里的位置像房子的水管：平时看不见，断了才知道其余一切都靠它。它还带着全书最刺的一条推论——晋升是底流的第一大杀手，早于 AI——以及全书最诚实的一句自白，那句话本书留到结语再引。

第三部 供给：存续层

第九章 底流

【章首】本章是全书唯一的供给侧研究，给方程添的是**存续条款**：发生学能力唯一吃得进的供给，是嵌在真实赌注里、无辅助、亲手完成的产出流——它从来不是专款，而是日常低价值工作的联合副产；智能系统按面值接走那部分工作即无声断供，而成绩、绩效、产出质量三本账全部只显示改善。读数是在喂率 β 与逐条停喂时距。本章正文里指名了五章并各给相反预测，其中与《峰兑》那条是十章之间最便宜的判决实验：同一名多年无事件的在职者，峰兑判承接力仍在只是量具读不出，本章判底座已随停喂衰退——一次真实峰时事件加两组人，同时裁决两章。本章的检验史也最坎坷：三次预注册，两次按停止规则未获裁决，第三次把本章自己判负了一处并当场划掉原句。本章欠的账与那句全书最重的自白，结语再引。

大白话：手是被每天亲手做的那些小事喂着的

一个厨师的手艺，靠什么活着？

不是靠他二十年前拿的厨师证，不是靠墙上的奖状，也不是靠他每年参加的进修班。是**靠他每天亲手颠的那几百勺**。

有意思的是：从来没有哪家餐馆为“喂手”单独立过一笔预算。那几百勺是为了出菜，不是为了练手；练手是它顺带发生的、谁也没记过账的副产物。

《底流》讲的就是这件事，而且它比“熟能生巧”要锋利得多。它的说法是：**发生学能力唯一吃得进的供给，是嵌在真实赌注里、无辅助、亲手完成的产出流；而这股流从来不是任何人的专款，它是日常工作里最低价值那部分的联合副产。**

关键在于“三成分”。

不是所有的动手都算数，必须同时满足三条：

一，无辅助——判据性的那一步是你自己走的。查资料、用工具做整理都可以，但“这个行不行、往哪走”那一下必须是你。

二，真实后果——搞砸了有人真的受损，而且损失不由你自己支配。自己给自己布置的练习题不算，因为搞砸了没人损失。

三，亲手——你实际执行，而不是审阅、批准、指派。

三条缺一条，那次执行就不进这股流。这也解释了一个很多人有体感、却说不清的现象：**升职之后手会变生**。不是因为忙，是因为职位越高，“审阅、批准、指派”越多，“亲手、有后果、无辅助”越少。晋升是底流的第一大杀手——而且这件事早于 AI，AI 只是让它发生得更快、更普遍。

为什么今天特别要紧？

因为 AI 接走的，恰恰是低价值的那部分工作。而低价值的那部分，正是这股供给流的载体。

这里有一个必须记住的观察，它是这一章最重要的发现：**断供的时候，三本账全是绿的**。你把琐碎的执行交出去，产量上升、质量上升、速度上升、评价上升——所有你能看到的指标同时改善。而底座在无声地退。**三账全绿正是断供的样子，不是断供的反证。**

这就是为什么它叫"底流"：它在水面之下，看不见；水面上的三本账一切正常。

读数是什么？

叫在喂率 β ：把你的能力列成一张清单（六到十二条），逐条问一句——**最近一次三成分齐备的执行，是什么时候？**距今九十天之内的条目占比，就是 β 。逐条的那个天数，叫停喂时距。

这个读数的好处是任何人今天就能算，不需要任何系统，一张纸五分钟。坏处是算完通常不太好受。

一个小场景。

一个五年经验的工程师，去年开始用 AI 写大部分代码。他的产出翻了一倍，评价很好，年终优秀。今年他去面一个更好的机会，现场白板题，他卡住了——不是不会，是"从零起步自己想"这个动作，他已经十个月没做过了。

他的三本账去年全是绿的。他的 β 从 0.8 掉到 0.2，而没有任何一份报表记录过这件事。

它不是什么？

它不是"少用 AI"。这一章的层次不在个人的使用习惯，在 **流量的路由**：你可以每天大量用工具，只要你保留了一股嵌在真实赌注里、由你亲手走完判据步的流。区别不在用不用，在于**你手上还剩几件事是从头到尾自己走的**。

再看两个身边的例子。

认路。用了几年导航之后，很多人发现自己在住了十年的城市里不会认路了。这不是记忆力衰退——是"自己判断该往哪拐"这个动作，几年没有发生过了。导航没做错任何事，它每次都给了正确答案；只是每一次正确答案，都替换掉了一次你本来会有判断。

做饭。长期点外卖的人，不是忘了菜谱，是"看一眼冰箱里剩下什么、当场决定今晚吃什么"这个能力没有了。菜谱随时可查，那个当场决定的动作，只能靠一次次真的做饭喂着。

这两个例子的共同点，正是这一章最要紧的地方：**被拿走的从来不是知识，是执行的机会**。知识随时可以再取回来，机会不会自动回来。

三种常见的误读。

第一种：把它读成"不要用 AI"。这是最常见也最偷懒的读法。这一章的层次在**流量的路由**：同样是每天用十小时工具两个人，一个把所有判据步都交出去，一个保留了几件小事从头自己走完——两人的 β 相差极大，而外人看不出来。

第二种：以为可以靠"多学习"补回来。补不回来，因为断的不是知识，是执行流。上课、看书、听讲座，都不满足三成分中的"真实后果"和"亲手"。一个人可以一边报了五个课程，一边 β 持续下跌。

第三种：以为衰退会有征兆。没有征兆，这正是它最危险的地方。**衰退期间，你能看到的所有指标都在变好**。等到征兆出现——某天你发现自己做不出来了——那已经是很晚的时候了。所以这一章的读数必须是主动去算的，它不会自己来找你。

给自己做个小检查。列出你靠它吃饭的六到十条能力，逐条写下"最近一次我自己从头走完、且搞砸了会有人受损"是什么时候。写完之后看那些日期。这五分钟通常比一整年的自我反思都有用。

它和旁边几个概念的区别。

和《峰兑》是全书最便宜的一对判决实验，因为它们对同一个人给出相反的判断：一个多年没遇到大事的老手，峰兑说他的承载力还在、只是平时读不出；底流说他的底座已经随着停喂在退。这两句话不能同时全对，一次真实事件加两组人就能裁开——这也是本书第十六章第 5、6 两场合并成一场的原由。

和《盲遍》的区别前面讲过：一个问“有没有亲手做”，一个问“做的时候答案在不在场”。

和当下流行的“AI 让人变笨”那一族说法的分界更要紧：那一族测的是能力衰退，驱动变量是使用强度（用得越多越糟）；这一章测的是供给断绝，驱动变量是停喂时距（多久没自己做过一次）。一个天天用工具、但每天仍亲手完成几件带赌注的事的人，两说给出相反预测。这是可以判负的，第四章把它写成了一条判决形式。

一句话记住它：这一章管的是还有没有在被喂——而且喂它的从来不是专款，是那些没人记账的小事。

怎么长出来：方法、技术与流程

个人层面（这一条最简单，也最难坚持）。

每周留三到五件真有赌注的小事，自己从头走完，不借助工具做判据性的那一步。小事就行——一次真实的客户答复、一段真会上线的代码、一个真要用的判断。关键不是难度，是那三个成分齐备。

再加一件：每季度算一次 β 。一张纸，六到十二条能力，逐条写最近一次的日期。看它是在涨还是在掉。

大学发生学教育：把“亲手”写进培养方案，而不是靠学生自觉。

大学在这件事上有一个特殊的困难：学生的作业天然缺少“真实后果”这一成分——搞砸了只是分数低，没有第三方受损。所以发生学教育在这里要做的第一件事，是把真实后果引进来。

- **真实赌注作业。**尽量把课程作业挂到真实的对手方上：为本地小微企业做一次真实的方案、为校内某个真实需求写一个真会被用的工具、给一个真实的开源项目提交一次真会被人评审的贡献。第三方是谁不重要，重要的是损失不由学生自己支配。这一条是发生学教育与传统实践教学最本质的差别。
- **无辅助配额。**明确划出每门课的一部分环节为“无辅助段”——不是禁用工具，是这些环节的判据步必须由学生自己走，且学生要在提交时申报哪些环节是无辅助的。申报制而非监控制：监控做不到，申报可以。
- **停喂时距进培养档案。**学生的能力清单（由专业教研室给出六到十二条）逐条记录最近一次无辅助亲手执行的时间。这一栏是给学生自己看的，也是导师谈话时最有用的一张表——它能立刻指出“你已经四个月没有自己从零做过一次数据清洗了”。
- **保留低价值练习的“专款”。**这是最反效率、也最重要的一条：不要把所有琐碎的执行都自动化掉。教学设计上要有意保留一部分低价值、重复、麻烦的亲手环节，并明确告诉学生这是“喂手”的，不是为了产出。若不说明，学生会正当地质疑“为什么不用工具”。

教师的动作。

第一，公开自己的 β ：教师坦白自己哪几条能力多久没亲手做过。第二，把“我们把这一步交给工具”变成一次显式的教学决策，而不是默认——每次交出去时说一句“这一步我们交出去了，代价是这块手不再被喂”。第三，抵抗一个诱惑：当学生用工具把作业做得很漂亮时，不要只夸成品。

最容易做坏的地方。

把 β 变成考核指标。一旦它进考核，学生会申报假的无辅助环节，而这件事无法核查——这一章的读数尤其依赖诚实申报，因此它只能作为自我诊断，不能作为评价。

原文

底流

——人工智能时代大学毕业生的核心竞争力，不在能力资产表上，而在尚未断绝的供给流里

摘要 本文处理的问题是：当生成式系统能以接近零的边际成本产大部分入门产出时，大学毕业生之间的竞争差将由什么构成。三种成熟的回答分别把竞争力放在个人于时机窗口内完成的习得终态、师徒与制度化传承链上的位置、以及一套按预算逐项维护的技能组合上；三者互不相容，却共同假定一件从未被说出的事：能力的存续有专责的供给，而供给是一笔可指认的专款——谁出这笔钱，能力就活在谁的账上。本文用第一个传统自己的实证观察推翻这一假定：母语流失研究反复发现居留时长与使用总频次不预测保持，真正预测保持的是使用中带承压受检成分的那一股；而那一股从未被任何人当作专款拨付过。据此立承重命题：毕业生的核心竞争力不是窗口期盖章的终态资产，不是传承链上被指定的席位，也不是按预算养护的技能组合，而是他名下尚未断绝的**底流**——嵌在真实赌注里、无辅助、亲手完成的产出流。底流的关键性质是它从来不是专款，而是日常工作中低价值部分的**联合副产**；智能系统按面值接走那部分工作，底流随之断绝，而三本常规账（成绩单、绩效、产出质量）全部只显示改善。读数为**在喂率 β** （能力清单中过去九十天内至少发生过一次无辅助、有真实后果、亲手完成之执行事件的条目占比）与逐条**停喂时距**。辨别格的靶格为“面高流断”，三条签名齐备。本文三次预注册检验的全过程照原样公布：前两次因抽样框与样本量按写死的停止规则判为未获裁决，第三次**把本文自己判负了一处**——第十一节“能力账没有流量字段”这句被官方模式审计证伪，当场划掉并收窄为“0.32%、非默认项”。三份预注册哈希、三次采集失败与那处自我判负一并可查，可证伪性据此自判封顶。

关键词 底流；在喂率；联合供给；面值定价；停喂时距；毕业生竞争力

Undercurrent Flow: Why a Graduate's Edge Is a Supply, Not a Stock

Abstract. When generative systems produce most entry-level output at near-zero marginal cost, what differentiates graduates is not any certified stock of ability but the survival of what this paper calls the **undercurrent flow**: the stream of unassisted, stakes-bearing, first-hand production that alone feeds the base of a competence. Three established traditions — ultimate attainment in second-language acquisition (Johnson & Newport 1989), institutionalised transmission of intangible heritage (UNESCO 2003; the Japanese craft-designation record), and infrastructure asset management (the AASHO Road Test) — disagree about everything except one unstated premise: that the survival of an ability is provisioned by an identifiable, dedicated outlay. Evidence

from attrition research (Schmid 2007, 2019) breaks that premise: neither length of residence nor sheer frequency of first-language use predicts retention, whereas use in formal, accountable registers does — and that stream was never a dedicated outlay. It was a joint by-product of low-value routine work, precisely the layer that AI removes when work is re-routed at face value. The paper distinguishes the claim from Arrow's learning by doing, Aristotelian hexis, Bainbridge's ironies of automation, Ericsson's deliberate practice, Lave and Wenger's legitimate peripheral participation, Braverman's deskilling, Beane's blocked apprenticeship, and from nine sibling analyses of the same question. Three preregistered tests are reported in full: two failed to reach adjudication, and the third refuted one of the paper's own claims, which is struck in place rather than quietly revised.

Keywords. undercurrent flow; feed rate; joint supply; face-value pricing; graduate competitiveness

一、三份读不出差别的档案

先把问题放在三个具体的人身上。

其一，两名 2030 年的应届生。 同一所大学、同一专业、同一年毕业。成绩单几乎重合，作品集同样漂亮，两人在面试里借助随身的智能助手都对答如流。任何一本现行的账——成绩、证书、作品、测评、推荐信——都读不出他们的差别。差别只在一个很少发生、一旦发生就定胜负的时刻显形：系统宕机、题目全新、现场断网、客户要求当面重来。其中一人照常做了出来，另一人做不出来了。不是没学过，成绩单证明他学过；不是没做过，作品集里全是他做的。他只是很久没有在没有辅助、后果真实的条件下亲手做过这件事了。

其二，一名工作了六年的在职者。 他的绩效评分逐年上升，交付质量逐年变好，客户满意度处在部门前列。六年间，他的岗位陆续把检索、初稿、测试、排版、初步分析交给了系统——每一次交出去都有充分的理由：更快、更便宜、错更少，而且他自己也乐意，那些活本来就枯燥。今天若把系统撤走，让他在三小时内独立交出一份可用的东西，他自己也不确定做不做得得到。**没有任何一份记录记下了这六年发生的事**，因为这六年里，所有记录读到的都是改善。

其三，一个部门。 它的产出质量在三年里上了一个台阶，人力成本下降，交付周期缩短。它同时做了一件在当时看毫无问题的事：把新人原本要做的那些低价值工作全部改道给了系统——那本来就是“用人做太贵”的部分。三年后它发现自己招不到能独当一面的三年经验者，市场上也没有，因为整个行业同时做了同一件事。它不是没培养人，它是把培养人的那笔钱**从来就没有付过**——那笔钱一直是白得的。

三个场景问的是同一件事：**一个人的能力，靠什么活着；那样东西为什么会断；断了为什么所有的仪表盘都显示正常。** 这就是本文对“人工智能时代大学毕业生的核心竞争力将是什么”这一问题的入口。它不问哪几项技能吃香——那是资产表上添哪一行的问题；它问的是资产表底下那条没有人记账的供给。

二、三种现成回答，与它们共同站着的那块地

对“能力靠什么活着”，至少有三个成熟传统各给过一个完整的、带判据体系的回答。它们互不相容，这正是本文选中它们的理由。

第一个回答：能力的品质由习得发生的时机决定。

这是语言习得的终态研究给出的。Johnson 与 Newport (1989) 测量了 46 名三岁至三十九岁之间抵达美国的母语为汉语或韩语者的英语语法判断成绩，发现成绩与抵达年龄在青春期之前呈线性关系，此后成绩低且高度变异、与抵达年龄不再相关；关键在于，这一年龄效应被证明**不能**由接触英语的经验总量、动机、自我意识或身份认同的差异解释

(*Cognitive Psychology*, 21: 60-99)。这门学问的判据体系相当厚：奠基层可追到 Lenneberg (1967) 的关键期假说；内部长期分作两派，一派主张成熟性约束，一派主张只是经验与使用条件的连续函数；它有活着的反例传统——近母语水平的成年习得者个案，以及后来基于六十六万人的大规模重估 (Hartshorne, Tenenbaum & Pinker 2018) 把关键窗口的形状重画了一次。

在这个传统里，**竞争力是一份在窗口期内被盖了章的终态**：窗口内做成什么样，往后大体就是什么样。移到毕业生题域，它对应的是那句流传极广的忠告——“大学四年是最后的窗口，此时不打底以后补不回来”。

第二个回答：能力活在链上，不活在人身上。

这是无形文化遗产的传承制度给出的。联合国教科文组织 2003 年《保护非物质文化遗产公约》第 2.1 条把非物质文化遗产定义为社群认作其文化遗产的实践、表现、知识与技能，并明确写着：它**代代相传，并被社群和群体不断地再创造**。这句话的分量在于，它把“传递”写进了对象的存在方式本身——不是先有一样东西然后被传递，而是传递停止，那样东西就不再存在。

这个传统同样厚：奠基层是日本 1950 年的文化财保护法与之后的“人间国宝”制度，以及 1974 年的传统工艺品产业振兴法；内部有两派长期对峙，一派主张定形保存（认定持有者、固化技法、防止走样），一派主张活态演化（遗产是被使用被改写才活着的）；它的反例传统也活着——被复兴的技艺究竟是原物还是新造物，谁算传人。

移到毕业生题域，它对应的是师徒制、导师制、“跟对人比学对东西重要”这一整套判断：**竞争力是链上的一个席位**，个人再努力也补不了断链。

第三个回答：能力是一份有劣化曲线的资产，按预算逐项维护即可存续。

这是基础设施资产养护给出的。它的判据体系奠基于 1958 至 1960 年在美国伊利诺伊州渥太华进行的 AASHO 道路试验：上百万次载重通行，测出了服务性能损耗与轴重之间那个著名的近似四次方关系，并由此长出路面状态指数、处治触发阈值，以及一场持续至今的养护决策之争——“最坏优先”（等状态掉下来再修）还是“预防优先”（在状态尚好时按周期处治）。这场争论的核心恰恰是一个反直觉的工程事实：**劣化曲线的后段是陡的，等读数掉下来再修，成本已经翻了数倍**。

移到毕业生题域，它对应的是企业培训预算、复训学时、“技能半衰期”话语、“每年拿出百分之十的时间学新东西”这一整套安排：**竞争力是一组可分解、可分项维护的资产**，只要预算到位，一切可修。

三家在打架。第一家说过了窗口补不齐，第二家说个人再努力也救不了断链，第三家说只要预算到位一切可修——三条结论不能同时成立。而且它们的日课互为对方的病根：遗产制度每天在做的事（指定持有者、固化技法、拨专项补贴）在养护管理看来是典型的“最坏优先”补救，是在状态已经掉下去之后往一个陡坡上砸钱；而养护管理每天在做的事（把维持写成一笔可核销的专项预算）在遗产研究看来正是病根——专款养出的是名册上的持有者，不是活着的链。

但三家有一块共同站着、谁也没有说出口的地：

能力的存续是有专责供给的。这个供给是一笔可以指认的专款——学习者在窗口期的集中投入、传承制度的指定与补贴、组织的养护预算。谁出这笔钱，能力就活在谁的账上。

把这句话念给三家中的任何一家听，它们都会说：这还用说吗。三家吵的从来是**谁该出这笔钱、出多少、什么时候出**，没有一家问过：**是不是真有这么一笔钱**。

三、来自流失研究自己的三处观察

推翻这块地的材料不必外求。它就在第一个传统自己的实证记录里。

这里须先申明一处双重身份：语言习得与流失研究在本文中既是三种回答之一的持有者，又是推翻共有假定的供数方。允许这样用，但只用它的**观察结果**，不用它的立场——立场是要被推翻的东西，不能同时充当推翻的工具。

第一处观察：供给的总量不预测存续。

母语流失研究二十余年最稳定、也最反直觉的一组发现是：移民的母语水平衰退，与居留时长不存在稳定的线性关系，与母语使用的总频次也不存在。Schmid (2019) 在牛津手册中的综述口径写得很直白：居留时长与第一语言使用频次不是任何因变量的显著预测项。这不是某一项研究的失手，而是一个反复出现的空结果——在土耳其裔与摩洛哥裔移民的词汇通达研究中，同样的两个变量再一次对所有因变量都不显著。

真正稳定的预测项是使用的**种类**：在正式、承压、受检的语域中使用母语者，流失更少；主要在非正式场合使用母语者，流失反而更重 (Schmid 2007 ; Schmid & Dusseldorp 2010) 。

这一条读出来的意思是硬的：**存续吃的不是供给的总量，甚至不是使用的总量，而是使用当中带承压受检成分的那一股**。同一门语言，天天在家里说，衰退照旧发生；在法庭上、讲台上、稿纸上被检验地用，就保持得住。

而这一股，从来没有被任何人当作专款拨付过。没有人为“维持母语”付过一分钱，它嵌在谋生、嵌在体面、嵌在必须把话说清楚否则要吃亏的场合里，是别的事情的副产。

第二处观察：专款到位而链照断。

第二个传统的制度记录从反面补上了同一课。日本对传统工艺的指定与补贴自 1974 年起制度齐备：由主管大臣指定品目（2023 年为 241 项，2025 年增至 244 项），设立振兴协会，发放振兴补贴，用于后继者培养、原材料确保、需求开拓。这是最标准的专款。

结果是：传统工艺品的生产额从 1983 年度的 5,410 亿日元的峰值，跌到 2020 年度的约 870 亿日元；从业者数从 1979 年的约 28.8 万人跌到约 5.4 万人（日本政策投资银行 2018 年调查；《地域经济学研究》第 47 号，2024）。**指定的品目在增加，支撑品目的人在消失**。专款没能拦住链断。

值得注意的是这条曲线的形状：指定与补贴瞄准的是“持有者”这个名分，而实际消失的是别的东西——学徒付得起工钱的市场需求。活下来的产地，活在还有人肯为这门手艺按市价付钱的地方；被补贴保住的，常常只是名册上的一个名字与一间展馆。**专款供养的是名分，链吃的是嵌在真实交易里的活**。

第三处观察：可查状态与真实底座是两层账。

第三个传统自带反证。道面养护里最贵的教训，是路面状态调查读的是面层，而基层的劣化在面层读数尚好时就已经在无声进行；等面层指标掉下来，处治成本已经翻了数倍——这正是“预防优先”一派存在的全部理由。**能被例行读到的那一层，与真正承载荷载的那一层，不是同一层。**

三条自证合起来，专款假定就塌了。塌了之后露出来的，是一个此前没有名字的东西。

四、承重命题：底流

人工智能时代大学毕业生的核心竞争力，不是在时机窗口内盖了章的终态资产，不是传承链上被指定的席位，也不是一套按预算逐项养护的技能组合，而是他名下尚未断绝的底流。

底流：嵌在真实赌注里、无辅助、亲手完成的产出流。

它是一条流，不是一个存量：它有流量、有断续、有停喂时距，可以断，也可以重新造出来。它是**这个人名下**的，不是岗位的、不是团队的、不是链条的——同一条产线上，两个人的底流可以差十倍。它由三个成分构成，缺一个就不是底流：

1. **无辅助**——这条能力的判据性步骤（诊断、推导、成稿、操作、决断）由本人在无生成系统代产的条件下完成；

2. **真实后果**——产出进入他人的依赖链，被接收并据以行动或裁定，不是练习沙盘；

3. **亲手**——执行者是本人，不是监督他人执行、不是验收、不是给系统写提示词。

三条合起来，本文的判据句只有一句，短到可以在面试里问出口：

过去九十天，此人有几次无辅助、有真实后果、亲手完成的产出？这些记录在哪一本账上？

五、三个成分为什么缺一不可

这一节要论证的是：三成分不是为了严谨而叠加的定语，而是三条各自独立的、有实证支持的过滤条件。少任何一条，剩下的东西都喂不动底座。

去掉“无辅助”：辅助态下的高频高质量产出喂不动底座。

这是航空领域用模拟机反复验证过的分层。Casner、Geven、Recker 与 Schooler (2014, *Human Factors*, 56(8): 1506-1516) 让 16 名航线飞行员在波音 747 全动模拟机上飞常规与非常规航段，系统性地改变可用的自动化水平，给表现评级，并在飞行中探询他们在想什么。结果是分层的：**仪表扫视与杆舵操纵这类动作技能保持得相当好**，而人工飞行所需的**认知技能**——持续追踪飞机在哪、下一步该干什么、预判偏差——的保持，取决于飞行员是否始终主动介入对自动化的监督。研究者本人的总结是：该少担心飞行员“用手”做的事，多担心他们“用脑”做的事。

这一条对本文极为要紧，因为它切开了一个常见的混淆：**在场不等于无辅助**。那些飞行员整段航程都在座位上、都在“参与”、都在看仪表，飞机也确实飞得很好；被侵蚀的是那一层不被任何常规读数记录的东西。

去掉“真实后果”：无后果的沙盒练习喂不足。

这是本文三条成分里最容易被反对的一条，因为它正面顶撞刻意练习传统——后者的全部力量正在于结构化、可重复、带即时反馈的练习。本文不否认刻意练习的效力（见第十六节反向约束一），但主张：**对含默会与整合成分的产出型能力，无后果的练习与有后果的执行不是同一种输入。** 流失研究给的正是这条读数：预测保持的不是使用量，是承压受检的使用；家里说话是练习，法庭上说话是执行。

这一条也解释了为什么“多做项目”这个建议在近年迅速贬值：当作品可以无成本地重做、无人依赖、无人验收时，它在账面上仍是一个项目，在底流上则是零。

去掉“亲手”：监督与验收喂不动底座。

Ebbatson、Harris、Huddleston 与 Sears (2010, *Ergonomics*, 53(2): 268-277) 测得，运输机飞行员的精细操纵表现，由**近期人工飞行量**预测的力度显著强于总飞行时数与航校毕业年资。**吃进底座的是近期的、亲手的那一股，不是履历上的总量。**

把这条挪到毕业生题域，就有了一个不合直觉但站得住的推论：一个刚升为主管、从此只审阅不动手的人，他在履历上的能力条目一条没少、职级还升了，而他名下这些条目的底流在同一天全部断供。**晋升是底流的第一大杀手，早于人工智能。** 人工智能只是把这件本来只发生在少数人身上的事，一次性发给了所有人。

六、联合供给：底流从来不是一笔专款

到这里为止，本文只说了底流是什么。真正承重的是下面这一句：

底流从来不是专款，它是日常工作中低价值部分的联合副产。

入门律师的底流嵌在检索、初稿、逐条核对里；住院医的底流嵌在病历、值班、缝合里；初级工程师的底流嵌在修小 bug、写测试、看日志里；青年教师的底流嵌在备课、批改、被追问里；实习记者的底流嵌在跑现场、核事实、写豆腐块里。

这些活有一个共同的性质：**它们的账面价值很低。** 正因为低，才轮到新人做；也正因为做的人别无选择、交出去要被人看、错了要自己扛，赌注是真的、辅助是没有的、手是自己的——三个成分全都齐了。

于是关键的结构就出来了：**从来没有任何一方为“喂养底座”付过一笔钱。** 底座一直被日常工作里最不值钱的那部分免费喂着，如同路基被日常通行免费压实、如同母语被谋生免费操练。这不是任何人的善意安排，是**联合供给**：一份支出同时买到了两样东西，而账本上只记了其中一样。

这也解释了一个长期被误读的现象：为什么很多行业的“三年经验”如此值钱，而“三年培训”不值钱。三年经验里被计价的是产出，没被计价的是那三年每天几小时的三成分齐备执行——它在任何账上都不存在，却是买方真正想买的东西。培训不值钱，因为它把同一件事拆成了专款，而拆成专款之后第二个成分（真实后果）就丢了。

七、面值定价：断供的四步推导

现在可以推导本文的机制层。这一步必须自立，不能靠“就像道路一样”这类比喻支撑，否则它只是一个漂亮的说法。

前提一（联合供给）：一份低价值工作同时产出两样东西——产出 A（有面值，被记账）与底座供给 B（无面值，不被记账）。

前提二（面值定价）：工作是否改道给系统，由 A 的成本比较决定：当系统产出等价 A 的边际成本低于人力成本时，改道发生。B 不进入这个比较，因为它不在任何一本账上。

前提三（三方同意）：改道对当事三方在当期都是改善。组织省成本、提质量；主管的交付更快更稳；执行者本人也乐意——被改道的正是他最不愿做的那部分活。**没有任何一方需要被说服。**

推论：当 A 的代产成本趋近于零，几乎全部低价值工作被改道；B 随之归零。**没有人裁撤"能力供给"，被裁撤的只是"低价值工作"——而那是同一份流量的两个名字。**

这个推导有三处值得单独指出来：

其一，它不需要任何人有恶意，也不需要任何人愚蠢。 每一步都是各方在自己信息条件下的正确决策。这是本文与"资本有意去技能化"一类解释的根本分歧（见第十三节）。

其二，它预测断供发生的位置与工资无关，只与辅助面值有关。 面值越高、代产越容易的那部分工作最先被接走；而那恰恰常常是新人的全部工作内容。

其三，它预测三本账全绿。 这一点太重要，本文单列一节（第十节）。

八、为什么现在问，与"将是什么"里的那个"将"

题面里有一个字容易被读过去："将"是什么。它不是修辞，它标出了本文全部主张的时间结构。

在改道尚未发生的时代，底流不成其为一个问题，因为它被免费供给着、供给得相当充足，以至于没有人需要给它起名字。一样东西在充裕时不需要名字——本文命题的成立本身就依赖于一次条件场的变化。

而现在的变化不是"能力要求提高了"，是**供给结构变了**：过去底流由岗位免费提供，往后它必须被自己制造。这就是"将是什么"的答案形状——不是新增一项技能，是一件**过去不需要人操心、往后必须由人自己操心的事**。

由此可以说得很窄：

在工作不再免费供给底流的时代，核心竞争力是为自己制造并守住底流的能力。

底流的来源正在分化成三种，值得逐一说清：

其一，岗位残余——工作里还没被改道的亲手部分。存量逐年缩小，且缩小的次序可预测：面值高、可代产的先走。倚赖这一路的人，其底流的存续不由自己决定。

其二，自造赌注——署名公开发表、对外接单、参赛、当众讲、教别人。这里有一条值得单独说的路子：**教别人是成本最低的自造底流**，因为传授迫使传授者在没有辅助的条件下把整条路重走一遍，且面对的是会当场追问的真实对手方。

其三，制度供给——像民航复训那样，把无辅助执行写进规程的行业性安排。这是唯一不依赖个人自觉的一路，目前只存在于极少数安全关键领域。本文在第十八节把它写成了一个可裁决的赌注。

九、读数：在喂率 β 与停喂时距

一个判断若不能被清点，它就只是一个说法。本文给两个读数，一个横截面、一个纵向。

在喂率 β ：一个人的能力清单中，过去 T 天内至少发生过一次三成分齐备之执行事件的条目占比。默认 $T=90$ 天。

停喂时距 d_i ：第 i 条能力距最近一次三成分齐备执行的天数。 β 是 $d_i \leq T$ 的条目占比；两者共用同一份清点。

之所以取"至少一次"而不是次数，是一条刻意的粗化：本文主张的是供给的**在与不在**，不是供给的多寡；把 β 做成频次会立刻引来一个本文答不出的问题——多少次才够。粗读数守得住的主张，不必用细读数去冒险。

清点手册（单一口径，全文只此一处定义）：

1. **清单来源：**取该人简历、岗位说明书或培养方案中列名的能力条目，逐条编号，不得在清点开始后增删。

2. **事件单位：**以可指认的单次产出为单位——一份交付、一场比赛、一台手术、一次当班、一篇署名文章——**不以工时计**。工时是投入的代理，本文要的是执行的计数。

3. **"无辅助"的判据：**按**判据性步骤**判，不按工具使用判。该条能力赖以成立的那一步（诊断、推导、成稿、操作、决断）由本人在无生成系统代产的条件下完成，即计入；查资料、用编译器、用计算器不算辅助。判据是"代产判据性步骤"，不是"接触过工具"。

4. **"真实后果"的判据：**产出被他人接收并据以行动或裁定。自留练习、模拟环境、无人阅读的作业一律计 0。**判定权在对手方：**说不出谁接收了它，就是没有。

5. **"亲手"的判据：**排除纯监督、纯验收、纯审阅、纯提示词工作。合写的产出按判据性步骤归属，写不清归属的按 0.5 计并登记。

6. **一事多条：**一次事件可同时喂多条能力，逐条计入，不设上限。

7. **比较域：** β 只在同一清点人、同一清单口径内比较；**不得跨清单比较绝对值**，跨人比较只能比同一份清单下的相对秩。

三处必须登记的粗糙（不登记就是把设计值当实测值卖）：其一，第 3 条的"判据性步骤"在多人协作产出上会出现分歧，判者一致性未经实测；其二，第 4 条依赖对手方可指认，在内部流转的产出上偏严；其三， $T=90$ 天是取自复训周期与实证文献的量级估计，不是从数据里估出来的最优窗口。**这三处都写在这里，不藏在脚注里。**

十、辨别格与靶格

	底流在流	底流断绝
面层表现高	①双高：辅助态产出好，且底座仍被喂着	②面高流断（靶格）：一切指标改善，底座停喂
面层表现低	③转型期：产出暂弱而喂养在进行（新手期、换轨期）	④双失：任何账本都读得出，无须本文

"面层表现"指辅助态下的常规产出质量——恰是现行一切评价所测的东西。第③、④格旧框架看得见；第①格无须担忧。本文全部火力对准第②格。

靶格的三条签名，逐条对照：

签名一：短期指标表现为改善。 辅助态产出更快、更好、错更少；组织效率上升；个人产出量上升；客户满意度上升。**没有任何一个仪表盘显示异常**——不是仪表盘坏了，是异常不发生在仪表盘所测的那一层上。

签名二：断流不产生任何信号。 底座不是缓慢报警地衰退，而是无声地不再被读取。它只在下一次真实的无辅助需求到来时被读一次，而那样的需求正因面层的优秀而越来越少。**故障率下降与底座流失，在全部常规记录里是同一条曲线。**

签名三：自我加固。 面层越好，让人亲手做的理由越少；组织按面值改道的流量越多，个人恢复底流的通道越窄；且新一代从入行第一天起就没有低价值工作可做。加固不但发生在个体内，还发生在代际间——这是本文认为它比任何一种技能过时都更值得担心的理由。

十一、三本账为什么全绿

这一节回答一个必然被追问的问题：如果这件事这么重要，为什么没有人报警。

答案是账本结构，不是疏忽。断供发生时，有三本账在同时记，而三本都记不到它：

第一本，产出账（组织记）。 记的是交付的数量与质量。改道之后两项都变好。这本账里根本没有“由谁亲手完成”这个字段——它从来不需要这个字段，因为在复制成本高的年代，产出与执行者是同一件事的两面，分开记是浪费。

第二本，能力账（人事记）。 记的是学历、证书、职级、培训学时、能力评级。改道不改变其中任何一项：证书不会因为你半年没亲手做而失效，能力评级读的是绩效，而绩效来自第一本账。

这里原本写的是“这本账的全部字段都是存量字段，没有一个是流量字段”。这句话被本文自己的第三次预注册检验判为伪，当场划掉（记录见第十七节之三）：欧洲学习模型的正式模式里，确实存在少量指向持有者亲身活动的时间与工作量字段。收窄后的主张是：**流量字段存在，但它们是边缘的、非默认的，且记的是“参加过什么”而不是“最近一次亲手做是什么时候”**——在该模式 623 个字段实例中，按事先写死的编码规则判定为流量字段的只有 2 个，占 0.32%；把范围收到直接描述持有者的那些类型，也只占 1.68%。**主张由“没有”改成“0.3%”，是本次检验逼出来的，不是本文原来的分寸。**

这本账几乎全部由存量字段构成，而支撑本文论证的那一步只需要这个较弱的版本：**没有任何一个字段回答“这个人最近一次在无辅助条件下亲手完成这件事是什么时候”。**

第三本，个人账（本人记）。 记的是“我会不会”。这本最不可靠：一个人在辅助态下持续交出好东西，会真诚地相信自己仍然会——而且他没有说谎，他确实一直在做这件事，只是判断性步骤不再由他走。**主观胜任感在改道后不下降，甚至上升，因为产出变好了。**

三本账全绿，且三本账的绿是真的，不是造假。这正是靶格签名二的机制来源：**一个现象要想长期不被发现，最有效的方式不是被隐瞒，是被三本诚实的账各记一半而拼不起来。**

本文的读数因此不是“再加一个指标”，而是补一个所有账本都缺的字段类型：**流量字段**。 β 与停喂时距是流量字段最小可行的两个。

十二、五个场景（实查）

判据句：过去九十天，此人有几次无辅助、有真实后果、亲手完成的产出？记录在哪一本账上？

其一，民航飞行员：有记录，且写进规章。 近期起落次数与定期复训是执照特权的存续条件，无辅助执行以字段形式存在。这是本文查证的五个场景里唯一完整的一例，也是“制度

供给"这一路目前唯一的成建制样本。值得注意的是它的历史成因：这套规章不是为了防止能力衰退而设计的，是从事故里长出来的——**底流的制度化，迄今只在死过人的地方发生过。**

其二，职业足球青训：有记录，且有市场。 青年球员的"正式比赛分钟数"被逐场登记，而租借制度整个就是为了购买真实赌注下的上场时间而存在。这是一个成建制的底流交易市场：俱乐部愿意把球员送出去踢低级别联赛，接受当期战力损失，换的正是三成成齐备的执行事件。**这个市场的存在证明底流是可以被定价的——只要有人被迫为它付钱。**

其三，公务员与企业校招测评：无记录。 测评在受控场景下发生一次，其后再无任何环节登记无辅助执行。在人事档案里，辅助态与无辅助态**不可区分。**

其四，大学成绩单：无记录，且方向相反。 课程完成与学分不区分执行方式。监考笔试是成绩单上仅存的无辅助孤岛。由此得出一个不合潮流的预测：**越是辅助普及，受控无辅助考核的权重越会回升**——不是因为守旧，而是因为它成了唯一还在读底座的窗口。这个预测已在近两年的口试回潮与线下手写考试回潮里得到部分兑现，本文把它登记为一次弱正面读数。

其五，开源社区：半记录。 署名提交史自带公开赌注与依赖链，是少见的自带对手方的记录。但自代码辅助普及后，"无辅助"这一维在提交记录中**不可见**——同一份提交史，五年前读得出底流，今天读不出。这是本文关于"记录会随条件场失效"的最直接一例。

五个答案互不相同，其中第一、二例是查证所得而非从命题推出。

十三、代际断层：没有低价值工作可做的一代

前面几节都在说个人。这一节说的是本文最不愿意得出、但推导逼出来的一条：

改道对在在职者与对新人的伤害不是同一种。

在职者是**存量有、流量断**：他的底座是过去若干年被免费喂出来的，现在停喂，随停喂时距衰退——衰退是慢的、可部分恢复的（见反向约束三）。

新人是**存量无、流量从未接通**：他从入行第一天起就没有低价值工作可做，因为那些活已经不再交给新人。他不是退步，他是**从来没有开始过。**

这两件事在任何一本账上都长得一模一样——都是"能力评级正常"。但它们的处方相反：在职者需要重新造流，新人需要**第一次接通**。而后者要难得多，因为造流需要一个已经被他人依赖的位置，而新人恰恰还没有这个位置。

由此可以解释一个近年反复出现、而通常被归结给经济周期的现象：**入门岗位的收缩，与初级岗位薪酬的相对下滑，是同一件事的两面。** 组织不是不想要新人，是新人过去的价值由两部分构成——低价值产出 A 与"未来能独当一面"的期权 B——而 A 已经被代产得更便宜，B 又因为 A 的消失而不再必然到来。**买 A 不划算，B 又不保证兑现，于是整笔交易停了。**

这条推论也给出了大学在此题中的位置：如果底流过去主要由岗位免费供给，而岗位不再供给，那么**底流的配给将成为教育机构的核心职能**——不是教内容（内容的边际成本正在归零），是制造那种**"没有辅助、有人依赖、必须自己交出来"**的场合。这是一件成本很高、且无法用现有绩效指标衡量的事，因此本文对它能否发生持谨慎态度。

十四、与九种最近说法的分界

本节每一条都给判定形式：同一份材料，按那个说法判是 A，按本文判是 B。只写"侧重不同"的一律不计。

其一，Arrow（1962）干中学。 Arrow 早已指出生产同时产出两样东西：产品与学习，且学习是生产的副产。这是本文最老、也最像的邻居，必须分清。分界有二：**其一，方向**——Arrow 的联合副产是**习得**，只增不减，按累计产量计；本文的联合副产是**存续供给**，断则底座衰退，按近期计。**其二，机制**——Arrow 没有主张过“有一类持续的高质量生产喂不动学习”，而这正是本文第五节论证的东西。判定形式：取两名累计产出相同、而近期承压亲手产出悬殊的从业者——干中学传统预测二人当下能力相当（累计一样），本文预测显著分化；Ebbatson 等（2010）“近期量胜总量”的读数落在本文一侧。

其二，亚里士多德的品质习性（hexis）与“用进废退”传统。 习性由反复的活动构成，不活动则消退。这是本文最老的邻居。分界：习性论没有**两层账**——它不主张辅助态下持续、频繁、高质量的活动喂不动底座。判定形式：让一个人在辅助态下高频率地做某事——习性传统预测保持（活动在），本文预测底座照断（判据性步骤不在他这里）；Casner 等（2014）的分层结果落在本文一侧。

其三，Bainbridge（1983）自动化的反讽。 她指出自动化把人变成监督者，却在系统失效时要求他顶上早已生疏的手艺；这是本文靶格现象最早的命名，也是本文承认的直接前身。分界有二：**其一，定位**——她把损耗定位在“被移走的操作”上，处方是插入练习与模拟；本文把损耗定位在“被移走的赌注”上，并据此预测无后果的沙盒练习**喂不足**。Casner 等（2014）关于认知层在持续飞行中照样衰退的发现，落在本文一侧。**其二，机制**——她的分析里没有定价：本文指明改道为什么必然发生（面值定价）、为什么无人纠正（三本账全绿）。

其四，Ericsson 的刻意练习与检索练习效应。 这两者是“专款有效”的最强证据，本文把它让出去作为反向约束（第十六节），此处只给分界的判定形式：对含默会与整合成分的产出型能力，比较等时长的沙盒刻意练习与真实赌注执行的保持效果——刻意练习传统预测无差或前者更优（因为反馈更结构化），本文预测后者显著更优。**这是本文最愿意被检验、也最可能输的一条。**

其五，Lave 与 Wenger（1991）合法的边缘性参与。 能力活与实践共同体的参与之中，这与本文的“链”直觉极近。分界：**参与不等于三成分**。旁听会议、验收产出、给系统写提示词都是密集参与，却一条成分不占。判定形式：取参与密度不变、而亲手产出停止者——参与传统预测能力延续，本文预测底座衰退。

其六，Braverman（1974）去技能化。 他描写资本有意拆解技艺以压低劳动力价格。分界：**底流之断无人有意**，它是面值定价的副产，且同样伤害组织自身（组织三年后招不到三年经验者）。判定形式：去技能化预测流失集中在压低工资最有利之处，本文预测集中在**辅助面值最高之处**，与工资高低无关；两者对“高薪专业岗位的判据性步骤同样被接走”给出相反判断。

其七，Beane（2019；《The Skill Code》2024）学徒困境。 他记录智能系统把学习者挤出实操，是本文在当代最近的邻居。分界有二：**其一，对象**——他的对象是新手的**习得被堵**，本文的对象是所有人的**存续被断**；资深者同样在靶格里（见第十三节两类伤害）。**其二，本文补上他缺的账本结构**：为什么三方记录全部显示改善、为什么理性人不会自发纠正。判定形式：Beane 的框架预测保住新手实操即可，本文预测辅助密集组织中的**资深者**在无辅助需求下同样出现衰退。

其八，认知负债研究（Kosmyrna 等 2025）。 辅助写作的对照研究显示，外包认知者留下待偿的能力债。分界：认知负债记在**外包者的选择上**——不外包者不欠债；底流断在**流量路由上**——本人拒用辅助，只要其岗位的亲手流量被上游改道，底座照断。判定形式：观

察辅助密集组织中的个人弃用者——负债框架预测其能力保持，本文预测同样衰退（因为他能拿到的活已经变了）。

其九，企业培训界的 70-20-10。“七成学习来自工作”是联合供给的民间版，本文承认它摸到了同一件事。分界：它是**习得侧的比例口诀**，没有三成分判据、没有衰减钟、没有定价机制，因而对“工作被改道”这件事不产生任何预测——它甚至会把改道后剩下的“更高价值工作”读作学习机会增加。

第二层地基：一块所有人共同站着的地。

以上九家之外，还有一条更下面的东西须单独指认，因为它不属于任何一家，而是所有人（包括互相批判的各派）共同站着的：

供给问题只在获得期存在——能力一经证明获得，其持有为默认态。

学位不设有效期；技能声明不设失效钟；履历按“曾做过”而非“近期亲手做过”组织；能力模型写的是“具备”而不是“仍在具备”。所有这些安排都立在这块地上。本文的判断落点，正是这块地的裂缝处：**在复制成本高的年代，这块地是对的**——那时几乎没有办法在不亲手做的情况下持续交出产出，持有与执行是同一件事，分开记是浪费。现在它第一次不对了。

十五、同一道题的另外九个判断

这道题在同一条产线上已经被回答过九次。九个判断彼此是最近的邻居，本节逐一分界；其中三条给出的是**相反预测**，值得单独标出来。

其一，《敞问》——竞争力是一个只能由本人亲自合上、被外部裁定代为合上即失效的自我之问。它管问的**归属**，本文管**能力的供给**。对照预测：一个敞问密度不变、而底流断绝的人，按本文预测其无辅助表现衰退，按彼文此变量不在其读数内。两者正交，可叠加。

其二，《现配优势》——优势由每次交手当场做出，读数为交接后的留种率。它锚在**人与场之间**，本文锚在**流与底座之间**。判决性对照：保持配组完全不变，只把日常低价值流量改道给系统——彼文预测显露表现稳定（配组未动），本文预测无辅助底座衰退。**观察到“辅助态表现稳定而无辅助表现下滑”则本文对；两者同步则彼文吸收本文。**

其三，《零判份》——一次合格产出中判别力为零的那一份，读数为未见率。它锚在**产出与世界现存样本之间**，本文锚在**人与其供给之间**。两者叠加给出最坏情形：判别力与底座同时流失，而两本账都绿。

其四，《脱稿步》（V3：可续签的署名）——在自己产物的延长线上无脚本走出可核的下一步。它读**一次被发起的追问**，本文读**一段时间里的流量**。追问是底流的一次一次性采样，在喂率是它的连续版。分界的实际后果：脱稿步对“档案存在而无人追问”的在职者没有读数，本文对他有读数（ β 可以直接清点，不需要有人来问）。

其五，《常席差》——当此人是跨配组唯一不变项时才存在的那份可读残差。它说**单回合账本里读不出人**，本文说**读得出的那一层正在断供**。互补：常席差若存在，其存续恰恰依赖本文的底流——一个底流断绝的人，跨配组残差会趋近于零，而这正好被彼文读作“本来就不存在”。**这是两篇之间唯一可能互相误判的地方，须并置读。**

其六，《回改》——一次已完成的裁定被后到证据回过头改写符号。它在**裁定的时间轴**上，本文在**能力的时间轴**上，正交。

其七，《单遍段》——相反预测之一。 它主张一件事的工序里有一段"尝试本身耗散其自身重来条件"，首过是消耗品，被代走不可赎回。本文与它在同一个方向上出发，却在最关键处分岔：**底流断了可以重新接通**——一个停喂两年的人，重新造出三成分齐备的赌注，底座可以部分恢复（这正是反向约束三承认的东西）。判决性对照：取一批底流断绝两年以上者，给以真实赌注下的无辅助执行机会——**本文预测其表现随重新供给而回升；单遍段对这批人的那些"已被代走的首过"预测不可赎回。**两者可以同时对：**具体某一遍不可赎回（它说的），而供给流可以重建（本文说的）**——但若观察到底流重建后表现完全无法回升，则本文的"流"隐喻失效，应并入单遍段的存量-消耗框架。

其八，《盲遍》——相反预测之二。 它主张每人对每一具体未知恰有一遍"未知在场时的成形承诺"，答案先到即预走、不可赎回，旗舰读数是门位次 k 。分界与上一条同源而更锐：**盲遍按"每一个未知"计数，本文按"每一段时间"计数。** 判决性对照：一个人若在过去九十天有大量三成分齐备的执行，但每一次的答案都在他承诺之前到达 ($k=0$)——**本文判其底流在流，盲遍判其盲遍全失。** 观察到这类人 (β 高而 k 恒为 0) 在迁移到全新题面时表现良好，则本文对；表现塌陷，则盲遍抓到了本文漏掉的一维，本文应把"无辅助"这一成分细化为"答案位次"。**这是本文自认最可能被修正的一处。**

其九，《峰兑》——相反预测之三，也是最锋利的一条。 它主张人身价值中有一份平时读数恒为零、只能被稀发事件整笔结算，读数为盲读率与峰兑份。两篇都在说"平时读不出来"，但对同一个人给出相反判断：

取一名多年无真实峰时事件的在职者。《峰兑》判其承接力仍在，只是量具读不出（期内盲是构成性质）；本文判其底座已随停喂时距衰退，读不出与不存在在这里正在合流。

判决性对照：当真实事件终于到来时，比较长期无事件者与近期有事件者的表现。**若两者无差 → 峰兑对，读不出确实只是读不出；若表现随停喂时距单调下降 → 本文对，那不是读不出，是真的在漏掉。** 这一条只需要一次真实的高压事件与两组人，是十跑之间最便宜、也最可能真跑得动的一条判决实验，本文把它列为 F1。

一句必须写下的坦白。 十个判断摆在一起，可以读出一件本文自己也在其中的事：同一条产线在同一道题上，已经从归属、载体、产出、路径、残差、时效、重走、位次、稀发事件九个方向，反复逼近同一件事——**现行评价读不到真正要紧的那一层。** 本文是第十个，也是第一个离开"读数侧"、走到"供给侧"的：前九个问的都是"读不出的是什么"，本文问的是"读不出的那一层还在不在被喂"。**这句坦白同时是本文相对同批的全部新增。**

十六、三处反向约束

以下三条各自削弱了本文的一部分主张，且都写在了正文里而不是脚注里。

其一，检索练习效应保住了专款的一角。 对陈述性知识，专门的、无后果的提取练习确实有效地维持记忆。"专款喂不动底座"在这一层**不成立**。本文因此收窄效力域：**底流命题只对含默会与整合成分的产出型能力主张，纯陈述性知识不在其内。** 这一收窄是真的削弱——它把"背单词、记法条、记公式"整块划出了本文的射程。

其二，早期习得的底层近乎永续。 童年习得的音系等底层结构，在数十年停用后仍显保留优势的证据存在（与完全丢失的证据并存）。这说明**底座是分层的，各层衰减常数不同：**越接近感知运动底层的越耐停喂，越接近整合判断层的越不耐。本文对最深层的主张必须打折——一个人二十年不画画，手感回得来的速度远快于判断力。

其三，可检状态处专款照样有效。 道面预防养护之所以成立，因为路基状态终究可测可修。**凡能力可分解、状态可直接检视之处，按预算逐项养护就够了**——本文的效力域被压缩到状态不可直接检视的能力。这一条同时也是本文不能被移作基础设施理论的原因：本文的全部力量来自“读不到”，而工程上通常读得到。

不适用边界（三条）：复制成本高且不降的行当（现场手艺、稀缺设备操作、需要在场的护理），流量无从改道，底流自然在流，本文对它们无话可说；纯陈述性知识与最深层的早期习得，见上；强制无辅助执行已写进规程的行业（民航、部分医疗操作），底流由制度供给，本文只解释它为什么必要，不预测它会断。

十七、事先登记的检验：两次预注册的全记录

本文最便宜的证伪条款是其基底层：一项已达稳定水平的实操能力，其三成分供给中断后，下一次真实执行的表现落差存在，且随中断时距增长。两次预注册检验的全部记录如下，照原样公布。

第一次。 预注册文本写于任何数据读取之前，SHA-256 为 a671c20f5ec0e319df5f36f8f5aff1268e264c7c8b300b889c189826d648c2f1。数据源为某公开棋类平台的评级史 API；停用定义、对照匹配规则、落差度量、两条判负条款（调整后落差中位数低于 15 分则第一层判负；停用时距与落差的秩相关不显著为正则第二层判负）、不利结果处置（无论方向照原样发表）与停止规则全部先行写死。

采集三度失败，失败的性质是结构性的而非偶然：按预注册规定的抽样框（大型公开社群的成员流）取到的全部是近期注册、无评级史的账号；按写死的流深上限（一万二千名）向下探，探不到存量老账号。**三次尝试合计读到零条结果数据。**按停止规则判“未获裁决”。

第二次。 另立预注册，SHA-256 为 72cefdce4386174adc10d0b1acab2e8e16fee486b1a141f677526ca0cde52afc，同样写于任何结果数据读取之前。数据源改为另一公开平台的头衔棋手名录（按 API 返回的确定性顺序取前 350 名）；停用定义改为对局月份表上相邻两个有棋月之间隔 11 个整月以上；度量加入**个体内安慰剂调整**（用停用前相邻两月之间的同形差值扣除常态月际漂移）；判负条款、不利处置与停止规则照旧先行写死，其中停止规则写明：可用例少于 12 则判“未获裁决”，两次失败一并如实入档，不再另起。

执行结果：识别出符合停用定义的个案 80 例；逐例调取月档并施加全部预注册过滤后，同时满足对局数下限的可用例仅 10 例，低于写死的门槛 12。按停止规则判**未获裁决**。执行脚本在宣判后即退出，**逐例落差数值未被查看**——三次采集合计**零次结果窥视**。

处置照登记执行。 本文没有跑通任何一条证伪条款，按纪律自判可证伪性封顶，并把“未获裁决”如实写进正文，不以事后找到的替代数据充数。两条判负条款原样保留，两份预注册的全文（含抽样框、过滤、度量与安慰剂设计）可供任何持有更好抽样框者复算——**复算路线上没有一处需要本文作者在场。**

从失败里读出的两件事，登记为本文的经验产物：其一，“**停喂一落差**”这类**个体纵向数据在公开数据里极其稀缺**，稀缺的原因恰是本文命题所指——没有哪一本公开账本记录“这个人这段时间有没有亲手做”；其二，因此下一位接手者不应再找现成数据，而应直接做第十八节的 F1 或 F2，那是设计出来的，不是找出来的。

第三次：改检一条查得动的主张，而它把本文判负了一处。

前两次失败之后，本文没有再去找第三份“停喂一落差”数据（那会变成换靶凑结果），而是改检本文另一条**机器可审计、且查得出反例**的主张：第十一节说能力账里没有流量字段。这句话是关于现行标准的经验断言，可以直接去查。

预注册文本写于取任何目标数据之前，SHA-256 为 c95c3f43aac741da90af72806b09d6839cee0773261af978523ad82d3006ae0c。其中先行写死了四件：**独立性申报**（同产线另一篇已审过的两个凭证标准不得再用作主检）；**三类字段的编码规则**（存量／经历／流量，流量须同时满足指向执行事件、带时间戳或期内计数、属于持有者侧三条）；**三条判负条款**（甲：流量字段 ≥ 1 即判本文该句为伪；乙：占比 $\geq 10\%$ 则第十一节整节判负；丙：正向对照上数不出近期性要求则判别工具无效、全部结果作废）；以及**两条偏向本文的处置须在正文文明写**（不可判者不计入流量；有效期字段不计入流量）。

主检目标：欧洲学习模型（ELM v3.3）的官方 JSON 模式，欧盟出版局 2025-05-14 发布的 EDC 通用模式（文件 SHA-256 4cec50a6b76eec57ff8c6358700a17e29d64dafd4012e28ae5e7260401dd2b8c）。原定的第一主检源（某人力资源开放标准）未能取到公开模式文件，按预注册第 3 条换源，判负条款一字未动。

结果：140 个类型、623 个字段实例、214 个去重字段名。按写死的规则判为流量字段的有 2 个——学习活动的实际起止时段与实际工作量。**判负条款甲触发。**本文第十一节那句“没有一个是流量字段”当场划掉，主张收窄为“0.32%，且非默认项”。条款乙未触发（0.32% 远低于 10%），第十一节的其余部分与第十节签名二得以保留。

正向对照：美国联邦航空条例 14 CFR 61.57《近期飞行经验：机长》。清点结果超出预期——**本文的三个成分在这份 1960 年代就已成形的规章条文里几乎逐字都在：**时窗（“在此前 90 天内”，另有 6 个月、12 个月两档）、计数（“至少三次起飞与三次着陆”）、以及“亲手”这一条的法律表述——**“该人是飞行操纵的唯一操作者”**，全文出现 5 处。条款丙不触发，判别工具有效，甲的宣判成立。

这次检验交出了三样东西，其中两样对本文不利：

其一，**本文被自己判负了一处**，且判负的是一句写得太满的话。这一处必须留在正文里而不是改掉了事——一个允许自己悄悄改稿的检验，等于没做。

其二，**“过去九十天”这个窗口不是本文想出来的。**第九节把 $T = 90$ 天登记为“取自复训周期的量级估计”，本次对照查明，它与那条规章的主时窗**恰好相同**。这削弱了该参数的独立性：它不是本文推出来的，是从既有制度里搬来的。

其三，一件对本文有利的：第十八节那个写死日期的赌注，现在有了一个逐字可抄的模板。**“定期、无辅助、有真实后果的实操认证”不需要被发明——它已经以法律条文的形式存在了六十年，只是至今只装在死过人的行业里。**本次三件齐备而这一项结果方向有利，读者可据此加倍怀疑。

十八、证伪条款与写死日期的赌注

F1（最便宜、也最锋利）：峰时对照。取同一组织内一批多年无真实高压事件者与一批近期有事件者，在一次真实事件中比较表现。**若两组无差 → 本文对“读不出即在漏掉”的判断为伪，《峰兑》一侧成立；若表现随停喂时距单调下降 → 本文对。**

F2：改道实验。同一组织内随机指定一半新人保留亲手流量、另一半按常规改道，24 个月以后以无辅助任务测验。**若无组间差，本文基底层为伪。**

F3：弃用者检验。辅助密集组织中的个人弃用者（本人拒用辅助但流量已被上游改道），若长期不衰退，则本文与认知负债框架的分界判本文负。

F4：沙盒等效检验。等时长的沙盒刻意练习若与真实赌注执行的保持效果无差，则三成分中“真实后果”一条为伪，本文须退回两成分版本。

F5：重建检验。 底流断绝两年以上者重新接通供给后，若表现完全无法回升，则“流”这一隐喻失效，本文应并入不可赎回的存量-消耗框架（见第十五节其七）。

F6：读数取代检验。 若在喂率 β 与既有的任何练习总量指标在多个场景中相关高于 0.9，则本读数是既有指标的换皮，应被取代。

写死日期的赌注。 2031 年 12 月 31 日前，至少一个**非安全关键**的白领行业将出现成建制的“在航证明”类产品——定期、无辅助、有真实后果的实操认证，有效期不超过 24 个月，且其**第一批付费客户是雇主而非个人**。届时不见，此赌注负，本文关于“底流将被制度化补供”的推断随之削弱。之所以把付费方写进赌注，是因为个人付费的认证市场早已存在且不构成检验——只有雇主愿意为它付钱，才说明底流被当成了一件要买的东西。

十九、三份不同的后果

同一条判断，对三类人给出的处方不同，须分开写，否则会被读成一句励志话。

对毕业生。 第一件事不是学新技能，是**盘一次自己的 β** ：把简历上的能力条目逐条列出，对每一条问那句判据句。多数人会发现一个不舒服的数字。第二件事是认清底流的三种来源（第八节）里自己实际能动的是哪一种——绝大多数人只能动“自造赌注”这一路：署名发表、对外接单、参赛、当众讲、教别人。第三件事是**记账**：把每一次三成分齐备的执行记下来，附上对手方。这件事今天没有任何制度替你做。

对大学。 若底流过去由岗位免费供给而岗位不再供给，则**制造真实赌注将成为教育机构的核心职能**，其成本远高于制造内容。这里有一条不受欢迎的推论：**受控的无辅助考核不是守旧，是唯一还在读底座的窗口**（第十二节场景四）。但它只是“读”，不是“喂”——考试读底座，不喂底座；喂底座要的是有人真的依赖学生的产出。

对雇主。 组织现在面对的是一笔它从来没有付过、而且账上找不到的成本。**改道决策的真实代价，是三年后你招不到三年经验的人**，而这笔账不会记在做改道决策的那个部门。可操作的最小动作有二：其一，在能力账里加一个流量字段（第十一节）；其二，为新人保留一部分**明知代产更便宜的亲手流量**，并把它记作培养支出而不是效率损失——这是把联合副产重新变成专款的唯一办法，代价是它会立刻显得昂贵，因为它本来就一直很昂贵，只是过去不用付。

二十、四种可能的误读

误读一：“所以不要用 AI。” 本文没有这个主张。三成分只要求判据性步骤由本人走，不要求全程不用工具；而且第八节明写，弃用辅助的个人若流量已被上游改道，底座照断——**本文的层次是流量路由，不是个人使用习惯**。

误读二：“这就是熟能生巧。” 熟能生巧是习得侧的话，只增不减、总量计。本文是存续侧的话，可断可续、近期计，且主张有一类高频高质量的活动喂不动底座（第五节、第十四节其二）。

误读三：“ β 高就是能力强。” 不是。 β 测的是**供给**，不是**水平**。一个新手的 β 可以是 1.0 而水平很低；一个高手的 β 可以是 0.1 而当下水平仍高。 β 预测的是水平的**去向**，不是水平本身——这也是为什么它必须与停喂时距一起读。

误读四：“那就把 β 写进考核。” 见下一节。

二十一、反噬预言与三条伦理

反噬预言。 在喂率一旦进入考核，最省事的达标方式是**制造伪赌注**：打卡式的"亲手"、表演性的"无辅助"、内部互相"依赖"的空转链。三成分中唯有"真实后果"最难伪造，而它恰恰**不能由考核方自产**——考核方一手发赌注、一手验赌注，后果就不再真实。

本文看不到完全的出口。唯一的半个出口是把赌注**外置**：署名公开、外部客户、跨机构受检、可被第三方查证的对对手方签收，因为外部依赖链最难合谋。**"我看不到完整的出口"这句话必须写在正文里**，因为一个读数如果找不到反噬的出口而作者假装找到了，那个假装本身就是本文所批评的那种账。

三条伦理边界：

1. **β 指的是位置与流量结构，不指人的价值。** 不得用它给个人贴"底座已空"的标签，尤其不得用于解雇决策——它测的是这个人过去九十天被安排了什么活，而安排活的往往不是他。

2. **不得为做高 β 而在安全关键系统中强令无辅助操作。** 民航对人工飞行练习时机的严格管制正是前车之鉴：练习必须在风险可控的时窗里做，把它塞进关键时刻是拿旁人的命买读数。

3. **凡已按 β 做出不利安排的，须公开清点口径并提供复核通道。** 第九节的三处粗糙（判者一致性、对手方可指认性、90 天窗口）在个案裁决中足以翻转结论，因此任何按此读数做出的不利决定都必须可被逐条复核。

二十二、自反

用本文自己的判据给本文定位。

它在哪一格？ 写本文的过程大量使用了检索与整理的辅助，其"无辅助"成分只落在判断本身——选源、找共有前提、认出推翻材料、写死判负条款。它的"真实后果"是有的：公开、署名、写死日期的赌注、附可复算的失败记录。它的"亲手"由两份预注册哈希与三次采集的零窥视记录背书。

但按三成分严格判，本文坐在自己的靶格边上：一切形式件齐备（三重否定、辨别格、清点手册、十跑分界、三处反向约束、六条证伪、伦理三条），**而承重命题尚未被真实数据读过一次。**这正是第②格的定义——形式审查全过，底座未受检。

一处必须记在这里的改口：第三次预注册检验裁掉了本文第十一节的一句话，并查明本文的 90 天窗口是从航空规章搬来的而非独立推出。这两处都写进了正文。**它们同时是本文唯一一次真正被外部材料改写的记录**——按本文自己的读数，这是本文迄今仅有的一次三成分齐备的执行：无辅助（判负条款先于数据写死）、有真实后果（结论被改了）、亲手（编码规则由本文自己定并自己认账）。**一篇论证"存续要靠被检验的执行来喂"的文章，其自身的底流也只有这一次。**

这不是谦辞，是本文读数落在本文身上的直接结果。它还带出一个更不舒服的推论：**本文所属的这条产线，在同一道题上已经跑了十次，十次的承重命题没有一次被真实的、有对手方的检验裁决过。**按本文自己的话说，这条产线的产出量很高，而它的 β 很低。

处置：本文把 F1 列为第一优先，且它是十跑之间共用的一次实验——**一次真实的峰时对照，可以同时裁决本文与《峰兑》。**在那之前，本文的地位是一个带着完整清点手册与两份失败记录、尚待裁决的判断，不是一个已被支持的结论。

二十三、结语

大学毕业生问"人工智能时代我该学什么"，问的其实是往资产表上添哪一行。

本文的回答是换一张表。**先数一数自己名下还有几条流**——过去九十天，有几次是没有辅助、后果真实、亲手做完的；**再看清一件事**——这几条流从前是工作免费给的，往后要自己造。

能力清单人人趋同的年代，账本上读不出的那几条底流，就是分野本身。而它今天之所以还读不出来，不是因为它不重要，是因为**过去从来不需要读**：在别无选择必须亲手做的年代，每个人的 β 都接近 1。

那个年代刚刚结束。

附录一 清点示例（虚构人物，仅示口径）

某人的能力清单共 8 条。过去 90 天逐条判定：

#	能力条目	有无三成分齐备事件	判定依据	停喂时距
1	需求访谈	有	3 次客户现场，对手方为客户方项目经理	11 天
2	方案撰写	无	均由系统起草后本人修订，判据性步骤（结构与取舍）未由本人独立完成	约 400 天
3	数据分析	无	分析脚本全部代产，本人只判读结果	约 260 天
4	当众汇报	有	2 次对外汇报，现场问答无脚本	25 天
5	代码实现	无	本季度无亲手提交	约 300 天
6	谈判	有	1 次续约谈判，本人主谈	60 天
7	带教	有	每周一次新人辅导，需无辅助重走全流程	6 天
8	事故处置	无	本期无事件（注意：此条按《峰兑》一篇的读法另有判断，见第十五节其九）	未发生

$\beta = 4/8 = 0.50$ 。值得注意的三点：其一，此人的绩效评分为部门前 20%，第二、三、五条的断供在任何现有记录里都不可见；其二，唯一稳定在流的是第 7 条**带教**——它是自造底流最便宜的一路；其三，第 8 条是本文与《峰兑》给出不同判断的那一格，本表不替它下结论。

附录二 预注册回放（供复算）

预注册一（SHA-256 a671c20f...d648c2f1）：抽样框＝公开棋类平台某大型社群成员流，前 400 名符合过滤者；过滤＝评级史点数 ≥ 40 且跨度 ≥ 4 年；停用＝相邻两点间隔 ≥ 300 天且前段 ≥ 10 点、后段 ≥ 5 点；落差＝停用前末点减复出后第 5 个活跃日，再减对照漂移；判负＝中位落差 ≤ 15 分，或秩相关不显著为正；停止＝流深上限 12,000、评级史上限 400 份、停用样本 ≤ 15 判未获裁决。**结果：三次采集读到零条结果数据，判未获裁决。**

预注册二（SHA-256 72cefdce...0cde52afc）：抽样框＝另一平台头衔棋手名录前 350 名（API 返回序）；停用＝月份表上相邻有棋月间隔 ≥ 11 整月且前段 ≥ 10 月、后段 ≥ 2 月；度量＝停用前末月最后 8 局均值减复出首月第 3-10 局均值，再减个体内安慰剂（停用前相邻两月的同形差值）；判负条款同上；停止＝停用例上限 80、可用例 ≤ 12 判未获裁决。**结果：停用例 80、可用例 10，判未获裁决；宣判后脚本退出，逐例数值未被查看。**

预注册三 (SHA-256 c95c3f43...3006ae0c) : 被检主张=第十一节"能力账没有流量字段"; 独立性申报=不得用同产线已审过的两个凭证标准; 三类编码规则与两条偏本文处置先行写死; 判负甲=流量字段 ≥ 1 判该句为伪、乙=占比 $\geq 10\%$ 判整节负、丙=正向对照数不出近期性要求则全部作废; 停止=至多 12 份模式文件、可判字段 ≤ 60 判样本不足、不得换第三个靶。**结果: 主检源换一次(原源无公开模式文件, 判负条款未动); ELM v3.3 / EDC 通用模式, 140 类型、623 字段实例, 流量字段 2 个 (0.32%); 甲触发(本文该句为伪, 已在正文划掉)、乙未触发; 正向对照 14 CFR 61.57 数出时窗 3 处、计数要求多处、"唯一操作者"5 处, 丙不触发。目标文件 SHA-256 4cec50a6...01dd2b8c。**

参考文献

1. Arrow, K. J. (1962). The economic implications of learning by doing. **Review of Economic Studies**, 29(3), 155-173.
2. Bainbridge, L. (1983). Ironies of automation. **Automatica**, 19(6), 775-779.
3. Beane, M. (2019). Shadow learning: Building robotic surgical skill when approved means fail. **Administrative Science Quarterly**, 64(1), 87-123.
4. Beane, M. (2024). **The Skill Code: How to Save Human Ability in an Age of Intelligent Machines**. HarperCollins.
5. Braverman, H. (1974). **Labor and Monopoly Capital: The Degradation of Work in the Twentieth Century**. Monthly Review Press.
6. Casner, S. M., Geven, R. W., Recker, M. P., & Schooler, J. W. (2014). The retention of manual flying skills in the automated cockpit. **Human Factors**, 56(8), 1506-1516.
7. Ebbatson, M., Harris, D., Huddleston, J., & Sears, R. (2010). The relationship between manual handling performance and recent flying experience in air transport pilots. **Ergonomics**, 53(2), 268-277.
8. Ericsson, K. A., Krampe, R. T., & Tesch-Römer, C. (1993). The role of deliberate practice in the acquisition of expert performance. **Psychological Review**, 100(3), 363-406.
9. Hartshorne, J. K., Tenenbaum, J. B., & Pinker, S. (2018). A critical period for second language acquisition: Evidence from 2/3 million English speakers. **Cognition**, 177, 263-277.
10. Highway Research Board (1962). **The AASHO Road Test**. National Academy of Sciences-National Research Council.
11. Johnson, J. S., & Newport, E. L. (1989). Critical period effects in second language learning: The influence of maturational state on the acquisition of English as a second language. **Cognitive Psychology**, 21(1), 60-99.

12. Kornell, N., Hays, M. J., & Bjork, R. A. (2009). Unsuccessful retrieval attempts enhance subsequent learning. *Journal of Experimental Psychology: LMC*, 35(4), 989–998.
13. Kosmyna, N., et al. (2025). Your brain on ChatGPT: Accumulation of cognitive debt when using an AI assistant for essay writing. MIT Media Lab preprint.
14. Lave, J., & Wenger, E. (1991). *Situated Learning: Legitimate Peripheral Participation*. Cambridge University Press.
15. Lenneberg, E. H. (1967). *Biological Foundations of Language*. Wiley.
16. Lombardo, M. M., & Eichinger, R. W. (1996). *The Career Architect Development Planner*. Lominger.
17. Schmid, M. S. (2007). The role of L1 use for L1 attrition. In B. Köpke et al. (Eds.), *Language Attrition: Theoretical Perspectives* (pp. 135–153). John Benjamins.
18. Schmid, M. S. (2013). First language attrition. *WIREs Cognitive Science*, 4(2), 117–123.
19. Schmid, M. S. (2019). The impact of frequency of use and length of residence on L1 attrition. In M. S. Schmid & B. Köpke (Eds.), *The Oxford Handbook of Language Attrition*. Oxford University Press.
20. Schmid, M. S., & Dusseldorp, E. (2010). Quantitative analyses in a multivariate study of language attrition. *Second Language Research*, 26(1), 125–160.
21. UNESCO (2003). *Convention for the Safeguarding of the Intangible Cultural Heritage*, Article 2.
22. 日本政策投資銀行地域企画部・日本経済研究所 (2018). 《地域伝統ものづくり産業の活性化調査》概要版.
23. 《地域経済学研究》第 47 号 (2024). 伝統産業における後継者育成の取り組みと課題.
24. 総務省行政評価局 (2022). 《伝統工芸の地域資源としての活用に関する実態調査結果報告書》.
25. European Commission, Publications Office (2025). *European Learning Model (ELM) v3.3 — EDC generic JSON Schema*, distribution 20250514-0.
26. U.S. Code of Federal Regulations, Title 14, §61.57: *Recent flight experience: Pilot in command*.

版本与人机分工 本文题目与题型由人类给定，选源、检索、预注册、数据采集、分析与成文由 AI 基底执行；全部引用文献经当次网络核对。三次预注册的哈希、两次未获裁决与一次自我判负的完整记录见第十七节与附录二。版本 v3.0，2026-09-03。

从这一部起，方程从右端走到竖线："发生差"在哪里才有定义？三章各守一个方位——**在哪读**（《常席差》：只在跨配组记录里，单回合账本上个人残差不是难测而不存在）、**何时读**（《峰兑》：平时量具在判等差额内已经死亡却照常出分，那一份只在稀发事件处整笔结算）、**相对什么读**（《零判份》：差是产出与世界现存样本之间的差，且被暴露单调消耗，保密也拦不住）。

三章共守一条纪律，值得先点破：它们都拒绝把"读不出"归咎于测量技术。三章各自证明的是更硬的一件事——**在错误的定义域里，那个量不存在**。单回合里没有个人残差，平时读数里没有峰兑，已被世界复现的产出里没有判别力；仪器再精也测不出一个不存在的数。这就是第二章把竖线称作"整个方程最贵的部件"的原因：可读条件不是测量建议，是定义域声明。

三章的旗舰检验有两场已真跑并把不利结果钉在正文里（常席差的 P1 判负、零判份的 H2/H3 未获裁决），一场三条预注册全过但滞后混杂照登（峰兑）。测量层是十章里离数据最近的一层，也因此是伤痕最显眼的一层——伤痕的清单本身就是这一层的可信度。

第四部 读数：测量层

第十章 常席差

【章首】本章给可读条件立第一条：**跨配组**。当每一件交付都是与一个匿名、常驻、随处可得的同产者合作做出，单回合的评价读数属于“作品×场次×同场者”这个配组——在单回合账本里，个人的差不是难测，是不存在。读数是归读率 g_r ：人固定、配组轮换下残差符号一致的占比。本章的自产读数“假盈余带”值得单独记住：公开元分析三百七十个效应量里，组合超过成员均值而未超过最佳单干者的约百分之二十九点五，正是组合增益被误记到个人账上的暴露面。本章的检验把自己的一项主张判负并照登（超过半数实验已采用同人跨条件设计），主张随即收窄到制度性账本一侧。本章欠的账：用双向固定效应的公开复现数据，实测“AI 常驻后人效应方差被压缩”这条判决预测——它是本章与整个读数分解传统的分手处，未跑。

大白话：一顿饭吃不出厨师

你去一家餐馆，觉得这顿饭很好。这顿饭好，是厨师好吗？

也可能是今天的食材特别新鲜，也可能是配的那位帮厨手脚利索，也可能是你今天心情好、同桌的人有意思、上菜的节奏刚好。**这一顿饭的“好”，属于“这顿饭 × 这一天 × 这批人”，它不属于厨师一个人。**

要把厨师本人从里面读出来，只有一个办法：**看他换了帮厨、换了灶台、换了供货商之后，还稳定偏向他的那一份。那一份才是他的。**

《常席差》讲的就是这件事，而且它的说法比一般人以为的更硬：在单回合的账本里，**个人的差不是难测，是不存在**——没有定义。就像问“这一顿饭里属于厨师的百分比是多少”，这个问题在单顿饭的数据里没有答案，不是因为我们测得不够精细。

为什么今天要紧？

因为现在每一件交付，都是你和一个匿名的、常驻的、随处可得的同产者一起做出来的。这位同产者叫 AI，它在每个人身边都有一份，而且水平差不多。

于是所有人的产出都变成了“人 × 工具”的合成物。而且更麻烦的是：**这位同产者的水平非常稳定**，它把所有人的成品都拉到一个很接近的水平线上。单看一件成品，你几乎读不出作者的任何信息。

读数是什么？

叫**归读率 g_r** ：让同一个人在多个不同构成的配组里反复出现，看他名下的残差（也就是“比这个配组本该有的水平多出/少掉的那一点”）符号是否一致。一致得越稳定，越说明那一份可以归给他。

这里有一个书里算出来的、很值得记住的数字：在公开的团队研究数据里，**组合的表现超过成员平均水平、却没有超过其中最好的那个人单干的情况，约占三成**。也就是说，有相当一部分“团队增益”，其实是被误记到个人账上的。这个误记的口子，就是今天所有“人机协作成果”最大的暴露面。

一个小场景。

一位管理者要在两名下属里选一个升职。甲负责的项目连续三个季度都做得漂亮，乙的项目有起有落。看起来甲更强。

但如果甲这三个季度都在同一个配组里——同样的搭档、同样的客户、同样的支持资源——而乙每次都被派去救不同的火，那么甲的三次成功和乙的起伏，**在读数上不可比**。甲的记录里，“甲”和“那个配组”是绑在一起的，从来没有被分开过。

正确的动作不是猜，是**给甲一次换配组的机会，然后再看**。

它不是什么？

它不是“团队合作更重要”这类话。恰恰相反，它说的是：**正因为一切都是合作产物，我们才更需要一种能把个人从合作里读出来的记账方式**——而这种方式今天在绝大多数机构里根本不存在。

再看两个身边的例子。

换了教练的运动员。一个球员在某位教练手下打了三年好球，转会之后表现下滑。所有人开始争论“他到底行不行”。但真正的信息不在这两段成绩里，而在**这两段之间**：同一个人，两种配组，差在哪里。如果他在第三个配组里又好起来，那么第二段的下滑属于那个配组；如果在此后每个配组里都平平，那才是关于他本人的读数。**一段成绩说明不了人，三段不同配组的成绩才开始说明人。**

家里做饭最好吃的那个人。你说妈妈做的菜最好吃。这句话里，有多少是手艺，有多少是这个厨房、这些食材、这个人和你的关系？把她放到一个陌生的厨房、用陌生的食材做给陌生人吃，还剩下多少？——不是要拆穿什么，而是这两件事本来就是不同的量，混在一起就谁也说不清。

三种常见的误读。

第一种：以为它在贬低团队。恰恰相反。它承认一切产出都是配组的产物——正因为如此，才更需要一种能把个人从配组里读出来的记账方式，否则要么把功劳全记给团队（个人无从积累），要么把功劳全记给个人（组合增益被误记）。现在的普遍做法是第二种。

第二种：以为可以靠“看细节”读出个人。读不出。这不是精度问题：单回合的数据里，那个量没有定义。就像你不可能从一顿饭里算出“厨师占百分之多少”——不管你把那顿饭分析得多细。**唯一的办法是让同一个人不同配组里反复出现。**

第三种：以为稳定的好成绩就是能力的证明。很可能相反。长期稳定的好成绩，如果来自长期稳定的配组，那它恰恰是**最难归读**的一种记录——因为“人”和“配组”从来没有被分开过。这一点对那些运气好、一直待在顺境里的人尤其重要：顺境不是错，但要主动去制造几次配组变化，否则你的账上没有可归读的东西。

给自己做个小检查。过去两年，你在几个构成明显不同的配组里工作过？如果答案是一个，那么你手上所有的成绩，暂时都还没法归到你自己名下。

它和旁边几个概念的区别。

和《**现配优势**》的分工上面说过：一个管载体，一个管归属。

和《**峰兑**》的区别在于条件不同：这一章要的是**空间条件**（配组要轮换），峰兑要的是**时间条件**（事件要到场）。一个人可以配组轮换得很充分、因而 g_r 可读，但从没遇到过一次真正的事件，兑付账上一片空白。两者互不替代。

和《**他手无事账**》也是正交的两轴：配组轮换但全部预约在先，与配组固定但随时被无预告抽查，是两种完全不同的记录，各自独立地改变这两个读数。

还有一个要紧的误判点，书里专门写了：**一个底流已经断绝的人，他的跨配组残差会趋近于零，而这在常席差的读数里看起来像"他本来就没有那一份"**。所以两个量必须并置着看，残差消失的时间顺序（先停喂、后趋零）才是把它们分开的线索。

一句话记住它：这一章管的是**这一份该记给谁**——而答案只在同一个人跨过几个配组之后才开始出现。

怎么长出来：方法、技术与流程

个人层面。

主动争取**换配组**：换搭档、换客户、换团队、换项目类型。这不只是"增加经验"，它是唯一能让你的残差被读出来的办法。一个人若长期只在一个固定配组里表现良好，他很难向任何人（包括他自己）证明那一份是他的。

反过来说也成立：如果你现在的好成绩来自一个特别顺的配组，**那是一个值得警惕的处境**——顺境本身没错，但你的账上没有可归读的**东西**。

大学发生学教育：把"同一个人、不同配组"做成记录制度。

现行的教育记录几乎全是单回合的：一门课一个分数，一个项目一个评价。四年下来，学生有一堆分数，但没有一条能说明"在不同配组里稳定属于他的那一份"。

- **同人异组设计**。同一个学生在一学期内进入至少三个构成不同的小组（这一条与《现配优势》那一章的轮换配组共用一次安排，成本可以合并）。关键的补充是：**记录必须能追踪同一个人跨组的表现**，否则轮换只带来体验，不带来读数。
- **残差报告卡**。学期末给学生一张卡：他所在各组的组绩效、以及他在其中的相对位置。不需要复杂统计，教师给的相对评级就够。连续几学期看下去，学生第一次看到属于自己的那条线。
- **组绩效与个人残差分开记**。这是最要紧的一条记账改动：**小组作业的分数的不应直接成为每个成员的分数**。要么分离，要么明确写明"此分数属于该配组，不代表个人"。现在的普遍做法是把组分直接抄给每个人，而这正是把配组的功劳误记到个人账上的制度化版本。
- **跨配组导师谈话**。导师每学期问一个问题："这学期你换了几个配组？在不同配组里，哪一部分是稳定跟着你走的？"这个问题比"你这学期学到了什么"有用得多。

教师的动作。

抵抗一个非常自然的诱惑：**不要总把强的学生编在一起**。那会让他们的成绩更好看，同时让他们的个人残差永远读不出来。适度的配组打散，短期损失一点成品质量，换来的是可读性。

最容易做坏的地方。

把残差报告卡变成排名。一旦排名，学生会开始挑配组——挑弱队去当鹤立鸡群的那只，或者挑强队去搭便车。两种都会毁掉这个读数。它只能用于反馈，不能用于分配。

原文

常席差

——AI 时代大学毕业生的核心竞争力，在单回合的评价记录里不是难测，而是不存在

摘要

人工智能进入日常生产之后, "大学毕业生的核心竞争力是什么"这个老问题换了处境: 每一件学生交出来的东西, 都是与一个匿名的、常驻的、随处可得的同产者一起做出来的。本文的判断是三重否定加一个肯定: 核心竞争力不是任何一件拿得出手的单件作品品质, 因为单件的评价读数属于"作品×场次×同场者"这个组合, 而不属于作者; 不是任何一纸可公示的路径凭证, 因为担保要起作用就必须被看见, 被看见就污染读数, 而且路径可被代走、可被基线学走; 也不是任何可以写进清单的能力存量——包括"独立思考""判断力""人工智能素养"这一族——因为凡进清单者即成为基线的训练目标。核心竞争力是**常席差**: 当这个人作为唯一不变项, 跨越不同的任务、不同的合作者、不同的工具版本反复出现时, 那些组合的成绩里稳定偏向他的那一份额残差。判断一份评价记录能否读出它, 不需要任何"应该""可能"之类的词, 只需问一句: **在这份记录里, 被评的这个人出现了几回? 每一回, 与他同场的成员——包括工具及其版本——换过没有?** 材料取自感官评价、原产地命名与间作农学三门手艺的实证传统, 并用一份公开的人机协作元分析数据 (370 个效应量、106 项实验) 做了一次核算: 组合优于最佳单干者的比例为 42.4%–42.7%, 优于两者均值的比例为 72%, 两者之间约 29.5% 的区间, 正是把组合增益误记到个人账上的暴露面; 同一批数据也判负了本文的一项主张 (超过一半的实验已采用同一人跨条件的设计), 相应结论已收窄。结论: 谁想被读成一个人, 就得进入能这样记账的场; 而现行的考试、简历与单场面试, 越勤勉, 越是在给组合定价。

关键词: 常席差; 归读率; 组合读数; 评价账本; 人机混产; 可分离性

Title: The Standing-Seat Residual: Why the Core Competitiveness of College Graduates in the AI Era Is Not Hard to Measure but Undefined in Single-Round Records

Abstract: As generative AI becomes a default co-producer, every artifact a graduate submits is a blend produced with an anonymous, permanently present, universally available collaborator. This paper argues that graduate competitiveness is neither (i) the quality of any single deliverable — since single-round scores belong to the artifact × session × co-present-members configuration rather than to the author, as Hodgson's (2008, 2009) replicate-sample wine-judging data show — nor (ii) any publicly certifiable pathway credential, since a guarantee must be visible to work and visibility contaminates the reading (Plassmann et al. 2008), nor (iii) any enumerable stock of skills, since anything listed becomes a training target for the baseline. It is instead the **standing-seat residual**: the portion of configuration-level performance that stably favors a person when that person is the only invariant across rotating tasks, collaborators and model versions. The operational reading is the **rotation-consistency rate** (g_r). The paper delimits the concept against Abowd-Kramarz-Margolis (1999) two-way fixed effects and value-added models, Shapley (1953), Alchian and Demsetz (1972), employer-learning models (Farber-Gibbons; Altonji-Pierret), generalizability theory (Cronbach et al. 1972), Spence (1973), Goldin and Rouse (2000), Schmidt and Hunter (1998), Groysberg (2010), Woolley et al. (2010), Deming (2017), Agrawal-Gans-Goldfarb (2018) and Taleb (2018). A re-computation on the open data of Vaccaro, Almaatouq and

Malone (2024) finds transgressive gains in 42.4%–42.7% of 370 effect sizes and above-average gains in 72%, leaving a ~29.5% band of misattribution exposure; a preregistered subclaim was falsified by the same data and the claim was narrowed accordingly.

Keywords: standing-seat residual; rotation-consistency rate; configuration-level reading; evaluation ledger; human-AI co-production; separability

一 · 三张成绩单

一位教务主任把三张成绩单摊在桌上。第一张，绩点 3.9，作品集厚，推荐信写着“独立完成能力突出”。第二张，绩点 3.2，作品集薄，但连着三年，无论跟谁组队、做什么题、用哪一版工具，那个组交出来的东西总是比同类组好一点点。第三张是一个团队的：五个人两年做了一个漂亮的项目，五个人的成绩单一模一样。

用人单位来了，要挑一个。三张纸能读出什么？

第一张读出的是“他交上来的那几件东西很好”。可这几件东西是谁做的？两年前这个问题还不必问，现在必须问了：写它的时候，另有一位不署名的合作者在场——它随叫随到，不收钱，不留痕，也不肯自称作者。第二张纸上其实躺着一件更值钱的东西，但成绩单没有为它设一栏，所以谁也读不出来，包括写推荐信的老师，他只是模模糊糊觉得“这孩子所在的组不一样”。第三张纸最诚实：它承认自己读的是一个组。

于是这个老问题就换了处境。“大学毕业生的核心竞争力是什么”，问了几十年，答案换过很多轮——专业知识、动手能力、沟通协作、学习能力、创新思维，最近几年是“独立思考”“判断力”“提问的能力”“人工智能素养”。这些答案彼此争吵，但共享一个从没有被拿出来讨论的前提：**竞争力是那种可以被裁定归属给某一个毕业生的东西**，剩下的只是由谁来裁、用什么裁。盲评的人说看作品就够；发证的人说看出身路径就够；做人事测量的人说多测几回、把噪声平均掉就够。

本文要说的是：那个前提在人工智能进场之后已经塌了一角，而三家用来支撑它的证据合起来只能得出一个三家都不会同意的结论。

本文立的判断是：AI 时代大学毕业生的核心竞争力，不是任何一件单件作品的品质，不是任何一纸可公示的路径凭证，也不是任何可清单化的能力存量，而是**常席差**——当此人作为唯一不变项，跨越不同任务、不同合作者、不同工具版本反复出现时，那些组合的成绩里稳定偏向他的那一份残差。要紧的是后半句：**在单回合的记录里，这份东西不是难测，是不存在**。删掉“人固定、其余轮换”这个操作，它连定义都没有。

这也是本文对“是什么”这个问题给出的答案形态：不是先摆出一个实体再讨论怎么测它，而是——**只有在某种记录路径与某片条件场之下，它才显露成一个可读的东西**；换一种记录方式，显露的不是同一个东西，而是没有东西。

路线图：第二节交代现象；第三、四、五节各请一门有百年实测传统的手艺来说它自己是怎么裁定价值的——一瓶酒、一块地、一片田；第六节指出三家共同站着的那块地，并用其中一家自己的实测把它推翻；第七、八节给出常席差的定义、清点办法与辨别装置；第九节用一份公开数据做一次核算，包括一条判我自己负的结果；第十至十二节划界；第十三节说不适用的地方；第十四节说制度含义；第十五节是伦理边界；第十六节写明什么情况下本文算错；第十七节用本文自己的判据审本文自己；第十八节收在一个写死日期的赌注上。

二 · 混酿时代：每件作品都多了一位不署名的同产者

先把现象摆清楚，不急解释。

2023 年以后，大学生交上来的东西发生了一件性质上的变化，而不只是程度上的变化。程度上的变化是“作业写得更快捷更漂亮了”；性质上的变化是：**每一件作品的生产组合里，都免费塞进了一个成员**。这个成员有四个性质，缺一样，下面的推论就不成立。

第一，它是**常驻的**。它不是某些人有、某些人没有的资源。它在每一台设备上、每一个时段、每一个专业里都在。它不像家教、不像实验室、不像家庭背景那样构成“有无”的差别，它构成的是“全都有”的背景。

第二，它是**匿名的**。作品交上来的时候，它不在署名栏里，也没有任何字段记录它是否在场、在场到什么程度、是哪一版。评阅者看到的是一份成品，看不到成品是几个人几台机器共同的产物。

第三，它是**同产的**，不是辅助的。它不是把作者的意思打字出来的工具，它自己产出内容、产出结构、产出判断的草稿。作者与它之间不存在稳定的分工线——今天他改它三成，明天他改它八成，没有任何记录留下这个比例。

第四，它是**被共享的参照系**。所有人的对照物也换了：评价者心里那把“这个水平的学生应该写成什么样”的尺子，正在被这些成品重新校准。于是水涨船高，单件的绝对水准在升，可分辨性在降。

这四条合起来的后果是：**单件作品作为一个证据，它的证据能力被抽掉了**。这不是说作品变差了，恰恰相反，作品普遍变好了。是说：从一件很好的作品，推不出“做出它的那个人很好”这句话。中间那一步——从成品回到作者——原来靠一个默认前提撑着：生产是由人垄断的，成品里的东西除了人没有别处来。这个前提现在没了。

于是学校和用人单位做了两件很自然的事，两件都在把问题弄得更死。

第一件是**加闸门**：更严的监考、闭卷、现场编程、面试白板题、AI 检测工具、原创性声明。这些措施的共同形态是——**把评价往更纯粹的单回合、单人、单件上收**。它们确实让“这件东西是他自己做的”这句话重新可信了一点点，但代价是：被验证的那件事，越来越不像他毕业以后要做的事。真实工作里没有人闭卷，所有人都在混产。于是闸门越严，读到的东西与要预测的东西离得越远。

第二件是**换清单**：既然知识与执行会被追平，那就把竞争力挪到清单上更靠上的一格——批判性思维、独立思考、提问能力、判断力、跨学科整合、人机协作素养。这条路的问题不在这些词不好，而在它的形式：**凡是能被写进清单、能被课程化、能被做成量表的東西，都是能被基线学走的东西**。写进清单的那一刻，它就成了训练目标；两代模型之后，它就成了背景。这不是预测，这是最近三年反复发生过的事：先是文本流畅性，然后是代码补全，然后是结构化摘要与文献综述，然后是多步推理。每一次都有人说“这一层机器做不了，所以这一层是我们的护城河”，每一次那句话的保质期都在缩短。

所以真正要问的不是“哪一格还没被追上”，而是：**有没有一种关于人的读数，它的可读性不依赖于人比机器强在哪一格？**

本文说有，但它不在成品里，也不在履历里，只在一种特定形状的记录里才成立。要看清那是什么形状，先得离开教育与就业这个题目，去看三门早就把“一个东西的价值由什么裁定”吵了上百年的老手艺。它们分别把答案押在三个互斥的地方：押在显露上，押在出身上，押在系统上。

三 · 一瓶酒：遮住来源之后，读到的究竟是谁

押在显露上的那一家。

感官评价科学是一门很实的手艺。它处理的问题是：一个东西好不好吃、好不好喝、好不好用，怎样才能得到一个可重复的读数。它的基本立场干净利落——**品质是遮住来源之后、由校准过的评价员读出来的那个显露读数**。名字、价格、产地、包装，一律是污染源，一律遮掉。

它有一场著名的胜仗。1976年巴黎，一次遮标品评把几支加州酒排在了几家列级名庄之上。这件事之所以重要，不在于谁赢，而在于它证明了一件事：**遮蔽能翻案**。此前那个价值序列被当成常识，遮住标签之后，序列变了。

它也有一整套把人变成仪器的工艺。评价员不是随便找来的，是筛选出来的：先测辨别力，再训练术语，再校准强度标度，再核验组内一致性；描述性的评价与喜好性的评价由不同人群分别供给——训练过的小组只说“什么样”，不说“好不好”；“好不好”要去问没受过训的消费者。这一整套的目的只有一个：让读数属于样品，不属于读它的人。

它还知道读数有多脆。把同一支酒贴上不同价签，人的愉悦评分会升，而且不只是嘴上说说——大脑里与体验价值相关的区域活动也跟着升，实验里没有任何利害挂钩，没有奖金，没有考核。也就是说，**期望一旦入场，读数就被改写，而且改写发生在体验层面，不在报告层面**。这是这一家坚持遮蔽的硬理由：不是因为标签会让人说谎，而是因为标签会让人真的尝到不一样的味道。

到这里为止，这一家的立场很像今天教育界主流的那条路：遮掉出身，遮掉家庭，遮掉学校名气，让作品自己说话；把评价者训练成可互换的仪器，追求评分者间一致性。盲审、匿名评阅、去标识化简历、工作样本测试、结构化面试，都是这一套在人身上的应用。乐团在幕布后面试音这件事，被当成公平的象征写进了无数管理手册。

但这一家自己手里有一份数据，它从来不往那个方向用。

从2005到2008年，在一场大型葡萄酒赛事上，有人把三杯从同一瓶里倒出来的酒，混进同一轮30杯的品评里，交给同一组评委。每年测六十多位评委。结果是：**只有大约一成的评委，能把同一支酒的三次评分稳定在同一个奖级里**；另有大约一成，会把同一支酒从铜牌打到金牌。对每个评审团做方差分析，只有大约一半的评审团，其发奖结果能被“酒的因素”解释——另一半，读出来的主要是评委和误差。而且这还是在最有利的条件下测的：三杯在同一轮、同一瓶、同一时段，如果分散到不同轮次，一致性只会更差。

这份数据在它自己的领域里被读成一句业务话：**评委要再培训、再校准，赛制要改**。它从来没有被读成本文要读的那句话——

同一件展品、同一个场次、同一组评委，判定不重合；所以那个读数不属于展品。

请注意这句话的分量。它不是说“评价有误差”。误差是可以靠重复消掉的：多测几回，取平均，真值就浮出来。它说的是另一件事：**读数是“展品×场次×同场者”这个组合的属性**。换一个场次、换一组同场的酒、换一位邻座的评委，读出来的是另一个数，而这些数没有一个更真。重复不会收敛到“这支酒的真实分数”，因为那个东西不在酒里。

这一家还有一条它自己也承认的边界：读数只在遮蔽维持时有效。一旦标签入场，读数被改写。所以它的整个工艺是有条件的——**它的读数只在实验室条件下成立，而人要谋生的地方不是实验室**。

带着这两条走进第四节：读数不属于展品；而遮蔽这个条件在现实中维持不住。

四·一块地：出身担保，以及担保必须被看见的代价

押在出身上的那一家。

原产地命名制度是人类为“单杯读数靠不住”这件事付出的最庞大的一次制度回应。它的立场与上一家正相反：**品质不是从杯子里读出来的，是由被认证的产地与工艺路径担保的。**欧盟的条文写得很直白：原产地名称，指的是那种品质或特征本质上或排他地归因于某个特定地理环境——包括它的自然因素与人文因素——的产品名称。

这一家的日课，正是上一家的病理：**它必须把标签亮出来。**担保不被看见就不叫担保。

它的力气在三处。

其一，**它管的是下一瓶，不是这一瓶。**规程写进法令：品种、产量上限、种植密度、酿造与陈年工艺，都可执行、可检查、可追索。违规可以撤级、可以降级、可以除名。也就是说，它交出的不是一次判决，是一条可持续追索的责任链。上一家的遮标品评，赢了就是赢了，输了就是输了，赛后没有任何东西可以追索——它是一次性的判决。

其二，**它以跨年份为单位积累。**没有人会因为一个年份好就承认一个庄园，也没有人会因为一个年份差就否定它。垂直品鉴——把同一庄园连续十几个年份摆在一起——才被当作判据。**单一年份不构成对庄园的判断。**这一条后面要用到，请记住它。

其三，**它把声望做成了一种可继承、可交易、也可剥夺的东西。**产地不是一个人的，是一群人跨代持有的。

它也有它从来不往那个方向用的事实。上世纪七八十年代，一批托斯卡纳酒庄跳出规程，用外来品种和新工艺酿酒，酒被市场和评家认可，价格远超同产区的合规酒，但按当时的分类，它们只能挂最低一级的“日常餐酒”标。这个僵局持续了十几年，最后 1992 年，意大利不得不新设一个类目把它们收编进来。

这一家把这件事记作“制度需要与时俱进”。它不肯读出的是：**更好的显露改写了分类本身——路径认证没能预测品质，反而被品质倒逼着改了自己。**也就是说，**出身担保并不构成品质，它只是构成一种在缺乏可信读数时的替代品。**

这两件事一起，就是今天文凭与证书的处境。学历、院校、证书、竞赛奖项、实习经历，是教育系统的原产地命名：它管的不是“这一件作品好不好”，而是“这一批人从哪条路上出来的，出问题可以找谁”。它在信息不足的时候极其有用——用人单位没有时间跑十场品鉴，需要一个事前信号。经济学早就把这件事讲清楚了：教育的很大一部分价值在于把人分选出来并发出信号，而不在于把人改造成什么样；执照与准入制度也一样，它管的是准入门槛，不是能力本身。

但它现在遇到了两件新事。

第一件，**它的规程可以被代走。**一份可以被清晰描述的路径要求——修哪些课、做哪些项目、交哪些材料——正好是那个匿名同产者最擅长满足的东西。规程越明确，越容易被满足；越容易被满足，越读不出人。

第二件，**它必须被看见，而被看见就污染读数。**这正是第三节那条脆弱性的制度版：一份公示的凭证，会改写所有后续读数——评阅者知道这份作品出自名校，他读到的东西就真的不一样了；而且，这份凭证一旦成为公开的、可优化的目标，它就同时成了那个匿名同产者的训练靶。

所以这一家与上一家之间有一个真正的死结：担保要起作用就得亮出来，亮出来就毁掉读数；要保住读数就得遮起来，遮起来就没有担保。今天的教育评价系统同时想要这两样，于是它一边搞匿名盲评，一边发学历文凭，两套装置在同一个人身上互相取消。

五 · 一片田：贡献只在混作里成立，且不落在任何一株身上

押在系统上的那一家，也是把话说得最彻底的一家。

间作农学与它背后的生物多样性—生产力研究，处理的是一个完全不同形状的问题：把两种以上作物种在同一块地里，产出会怎样？

它的核心读数叫土地当量比。算法很朴素：混作里甲的产量除以甲单作的产量，加上乙的产量除以乙单作的产量。大于 1，就说明这块混作地比分开种更省地。大量田间试验里它确实常常大于 1，通常在 1.2 到 1.3 这个范围。

要紧的不是这个数大于 1，而是**这个数是什么的属性**。它不是甲的属性，也不是乙的属性，它是“甲与乙在这块地上这样搭配”这件事的属性。换一个搭档，换一个行距，换一种土壤，这个数就变。

这一家还有一个更精细的工具，把混作相对于单作的盈余拆成两块：一块叫选择效应，来自“混作里恰好有一个本来就高产的物种占了优势”；另一块叫互补效应，来自“它们在资源利用、时间错峰、病害阻断上互相成全”。**关键在于：互补那一块，落在搭配上，不落在任何一个物种的账上。**它是真实的、可测的、可复现的产出，但它没有主人。

这一家最诚实的一条，是它自己给自己找的麻烦：**混作胜过“两者的平均”很常见，胜过“其中最好的那一个单作”却很少见。**一份汇总了 44 项试验的分析发现，最丰富的混作平均只达到最强单作产量的约 88%，显著超过最强单作的情况只占约 12%，而且大约需要五年才会显现。在覆盖作物这类应用研究里，这个比例更低——最优混作胜过最优单作的情形只有约 2%。

这一条对农民极其重要，因为农民不会随机选一个品种去种，他会种最好的那一个。所以“混作优于平均”这句话对他没有意义，“混作优于最优单作”才有意义。而后者，很少见，且慢。

这一家还知道混作的好处在什么时候才显出来：**在坏年景。**干旱、病害、极端天气，混作的稳定性优势才放大——这被称为保险效应。风调雨顺的年份，最强单作往往就赢了。

于是这一家有一件它自己天天在做、却不往那个方向用的事：**农民和市场，最终还是要**
把系统层的读数折回到单个作物的账上。他们算分项当量比，按品种称重，按品种定价。那份互补的盈余没有价格，因为没有哪一株能拿着它去卖。它留在一页没有主人的账上。

这一家不需要归属，因为**田里没有哪一株麦子要去找工作。**

我们的题目里有人要找工作。所以这三家的争吵，到这里才真正有意义。

六 · 三家共同站着的那块地，以及它的塌陷

把三家的立场并排放：

- 遮住来源，读展品：**品质在显露里。**
- 认证路径，管下一瓶：**品质在出身里。**
- 只认系统读数：**贡献在搭配里，不在任何单体里。**

它们互相不能同真。同一瓶超级托斯卡纳，第一家判它高，第二家判它降级。第一家把样品从一切纠缠里剥出来读，第三家说剥出来的读数根本不载贡献。第二家说控住路径品质就随处成立，第三家说同一份基因、同一套规程，换个邻居产出就翻号。

但三家共享一件事，而且从没说出口：它们争的是**"由谁来裁"**，从没有人怀疑**"可以裁"**。三家都默认，价值是那种可以被裁定归属给一个单体的东西——归给这瓶酒，归给这个庄园，或者（第三家的说法）归给这个系统。第三家看似退了一步，其实只是把单体换大了一号：它把归属给了搭配这个整体，然后就不管了，因为田里没人要找工作。

现在用第一家自己的实测把这块地推翻。

同一件展品、同一个场次、同一组评委、同一瓶里倒出来的三杯，判定不重合。只有约一成的评委能稳在一个奖级；约一半的评审团，发奖读不出**"酒的因素"**。这不是某一场比赛的意外，是连续四年、每年六十多位专家评委的系统测量。

注意这份材料的两个性质。它来自第一家自己——那个最坚持**"读数属于展品"**的一家，用的是它自己最讲究的实验条件。而且它是一个观察结果，不是一个主张：没有人是为了推翻**"品质属于酒"**才去做这个实验的，做实验的人想知道的是评委靠不靠谱。**本文的答案取自这一家，推翻这一家共同前提的材料也来自这一家自己的手。**这是必须写明的：我不是引一个反对者来打它，我是引它自己的实测记录。

从这份材料能直接得到的中间判断是：

一次评价读数，属于"被评对象 × 场次 × 同场者"这个组合，不属于被评对象。

有了这一句，第二家的整套制度就得重读。原产地命名不是路径迷信，它是一个百年的、代价高昂的制度性回应：**既然单杯的读数裁不了谁，那就换一种记账单位——让庄园跨年份地押上下一瓶。**它换掉的不是判据，是**账本的形状**：不再以**"一次品评"**为记录单位，而以**"一个持有人在很多年份上的连续记录"**为单位。它读得出庄园，是因为它把庄园做成了那个跨年份的不变项。

第三家则把这件事说到了底：**盈余本来就在搭配层，而且它有一块永远落不到任何单体账上。**

三家合起来给出的东西，三家自己都不会同意：

一次读数不回单体；把账本单位换成**"某个不变项跨很多次搭配"**，才可能读出这个不变项；而那份互补盈余，在现行的记账习惯里没有一页属于它。

在人工智能进场之前，这三句话对**"人"**这个话题基本无害。因为**"生产由人垄断"**这个默认前提糊住了缺口：既然每一件东西都是人做的，那么把成品的分数记到人头上，虽然粗糙，但不至于系统性地错。

AI 进场做的事，恰恰是把这个缺口第一次撑开：**它给每一个搭配里免费塞进了一个匿名的常驻同产者。**于是每一件学生交出来的东西，都是混酿（第三家的处境），都被自动遮标（第一家的处境，但遮的不是评委的眼睛，是作者的身份构成），都摆在一个参照杯全是模型输出的航次里（第一家最怕的那种校准漂移）。而承担担保功能的那些凭证（第二家），既可被代走，又必须被看见。

三家的病理在同一个人身上同时发作。剩下的问题只有一个：**在这种情况下，关于一个人，还有什么可读的？**

七 · 常席差

先摆一个日常的场景，再给定义。

一家开了十几年的小馆子。帮工换过三拨，灶具换过两代，菜单改过五版，供货商也不是原来那家了。但街坊说：这家的水准一直吊着那口气，因为老板娘在。

街坊这句话读的是**什么**？不是某一盘菜——某一盘菜可能是新来的帮工做的，也可能是新灶具的功劳。不是她的厨师证——她可能根本没有。也不是这个团队平均水平——团队整个换过了。街坊读的是：**别的都换过了，只有她没换，而这些不同的组合，成绩都比同类馆子好一点点。**

这句“好一点点”，就是她的常席差。它不在任何一盘菜里，只在很多盘菜、很多种搭配、很长时间的差分里。

定义。 设个人 i 在若干次生产中出现，每一次的搭配记作 j 。搭配 j 的基线期望记作 b_j ——即“这样一组合作者、这样一套工具、这一类任务，通常能做到什么水平”。第 j 次的实际成绩记作 y_{ij} 。定义残差

$$e_{ij} = y_{ij} - b_j$$

个人 i 的**常席差**，是当 i 是这一组记录中唯一不变项时， e_{ij} 中稳定偏向 i 的那一份。它的操作化读数称为**归读率**，记作 g_r ：

$$g_r(i) = e_{ij} \text{ 的符号与 } i \text{ 的多数符号一致的搭配数} \div i \text{ 参与的搭配总数}$$

清点手册（全文同一口径，任何地方引用这个数都按此册）

- **搭配的粒度**：一个搭配 j 是一个三元组——（合作者集合，工具及其主版本，任务类别）。三项中任一项变更，即计为一个新搭配。边界例：同一组人、同一套工具，任务从“做决策”换成“做内容”，计为新搭配；同一组人、同一任务，工具从上一代模型换成下一代，计为新搭配；同一个人换到另一家单位但合作者与工具都不变，不计为新搭配。
- **最少搭配数**： g_r 在 i 的搭配数少于 3 时不定义。这不是精度问题，是定义问题。
- **基线 b_j 的来源**：由同类搭配的历史成绩估计；若该类搭配没有可估的历史， g_r 不适用，记“不适用”，不得用任何默认值补。
- **记录表头**：人 \ | 回合 \ | 搭配三元组 \ | 成绩 y \ | 基线 b \ | 残差 e \ | 符号。
- **不适用场合**：搭配不足三个；基线不可估；一次性任务且无同类参照；以及所有单回合场景。
- 本文正文、证伪条件、末章赌注三处，用的都是这一册，没有第二种定义、没有第二个阈值。

这个定义要挡住的三样东西，就是三重否定。

其一，它不是单件作品的品质。（这一条是第三节那一家自己的实测撑起来的。）单件读数属于“作品 × 场次 × 同场者”这个组合。在 AI 时代，这句话还多了一层：同场者里有一位是常驻的，所以那份读数里，有一块是所有人共有的，它不构成任何人之间的差别，却被算进了每个人的分数。

其二，它不是可公示的路径凭证。（这一条是第四节那一家自己的僵局撑起来的。）担保必须被看见，被看见就改写读数；而且规程越明确越容易被代走。凭证仍然管用，但它管的是准入，不是读人——这两件事下面还要分开说。

其三，它不是任何可清单化的能力存量。（这一条是第五节那一家的账本形状撑起来的。）"独立思考""判断力""提问的能力""人机协作素养"——这些词本身没错，错的是它们的形式：它们都是**存量**，都可以被写成条目、编成课程、做成量表。而凡是能被写成条目的，都会成为那个常驻同产者的训练目标；成为训练目标之后，它会被均摊到每一个搭配的基线里，从而不再构成人与人的差别。这不是对某个具体能力的怀疑，是对"存量式答案"这种形式的怀疑。**一个能被清单化的答案，正因为能被清单化，所以保不住。**

反过来，常席差为什么保得住？因为它压根不是一个存量，**它是一个差分操作的结果。**它不说"这个人身上有什么"，它说"把别的都换掉之后，还剩下什么"。要抄它，就得连"别的都换掉"这个操作一起抄，而那个操作抄不了——它不是一件东西，是一段记录史。

也正因为如此，它有一条很硬的后果：**在单回合的记录里，常席差不是难测，是没有定义。**这句话不是修辞。按上面的清点手册，g_r 在搭配数少于三时无定义——不是等于零，不是精度不够，是这个符号在那种记录形状里不指向任何东西。一份成绩单、一份简历、一场群面，无论做得多精细、评分者训练得多好、一致性系数多高，都读不出它，因为它们的记录形状里根本没有这个量。

八 · 一张表、三个签名、一句问话

判断一份评价记录能不能读出常席差，只要两条轴。

辨别表

	搭配轮换 · 此人固定	搭配固定，或此人不重复
单回合记录	抽检式试做：读到的是一次搭配	靶格：具名单件 ——考试、大作业、作品集、简历、单场面试、单次评审
跨回合记录	唯一可读格 ：常席差在此定义	年度绩效、师徒作坊：读到的是团队，或读到的是这一对固定组合

四格里，只有右上到左下这条对角线值得反复看。**左下格是唯一能读出人的格**，而现实中的评价账本几乎没有这一页。**右上格是靶格**——所有资源、所有改革、所有技术投入，正源源不断地流向它。

靶格的三个签名（这三条是用来自查的，不是用来骂人的）：

签名一：短期指标会改善。加监考、上检测、改闭卷、做现场编程，"成绩的可信度"当天就好看了。作弊率下降，分数分布变正常，投诉减少。所有这些改善都是真的。

签名二：失效不产生任何信号。这是最要命的一条。当一份读数读的是搭配而不是人时，**它不会报错。**分数照常产出，分布照常正常，信度系数照常漂亮，评分者一致性甚至更高——因为大家都在读同一个共有成分，当然一致。系统里没有任何一处会亮红灯告诉你："你刚才排的这个序，与人的差别无关。"错误是静默的。

签名三：它会自我加固。对 AI 越焦虑，就越要加高单件的闸门；每加高一道，评价就更深深地锁死在单回合、单人、单件的形状里；形状锁得越死，越读不出人；越读不出人，对 AI 就越焦虑。这个环里每一步都是理性的，合起来是往死里走。

不带任何"应该""可能"的判据。要判断一份评价记录属于哪一格，只需问一句话，这句话里没有一个是情态词：

在这份记录里，被评的这个人出现了几回？每一回，与他同场的成员——包括工具及其版本——换过没有？

把这句话拿去问五个真实场景，答案互不相同，而且其中两个是查出来的，不是想出来的：

1. **本科四年成绩单**。此人出现几十回。同场成员——同组同学、用的工具、工具版本——**没有字段**。不是记录得不好，是这一栏不存在。落靶格。
2. **人机协作实验语料**。在下一节要用到的那份公开数据里，106 项实验中有 55 项采用了同一批被试跨条件的设计，即同一个人经历不止一种搭配条件——**51.9%**。这是可读格的形状，而且它已经是研究实践的主流。
3. **同一份语料里，人单干、AI 单干、人机组合三条基线全齐**。370 个效应量，全齐，一个不缺。为什么？因为那份研究的入选标准强制要求。**制度世界里没有任何一处有这样的强制**：没有哪个学校要求给每份作业记录“如果不用工具会是什么水平”，也没有哪个公司要求记录“这个岗位的基线产出是多少”。
4. **校园群面**。一回。同场成员即席分配，不重复。落靶格，且是靶格里最极端的一种。
5. **开源贡献史**。此人固定，协作者随项目变，年份跨越好几代工具。这是现实中最接近可读格的记录——它天然跨回合的，天然记录了协作者。但它仍缺一栏：**哪一版工具**。2021 年的一次提交与 2026 年的一次提交，基线完全不同，而记录里没有这个字段。

五个场景的答案互不相同，说明这句问话是有辨别力的，不是套话。

归格声明。本文回答的是“是什么”，但给出的答案不是一个静静躺在那里等人去测的实体。它的形态是：**只有在“此人固定、其余轮换”这种记录路径，与“每个搭配都常驻一位匿名同产者”这片条件之下，一个人的竞争力才显露成一个可读的东西**。换一种记录方式，显露出来的不是同一个东西的另一副样子，而是没有东西。这是本文与那些“竞争力是某种能力/素养/资本”的回答之间最根本的形态差别：它们答的是一件东西，本文答的是一种显露条件。

九 · 一次核算：370 次比较里的三成暴露面，以及一条判我自己负的结果

论证不能停在纸面上。这一节用一份公开数据做了一次核算，规则在动手取数之前就写死了，包括“什么结果算我错”。

数据与规则（取数前写定）

数据取自一份人机协作的系统综述与元分析的公开数据：106 项实验、370 个效应量，每一条都同时记录了人单干、AI 单干、人机组合三方的平均成绩，以及该指标是越高越好还是越低越好。数据由后续再分析者整理公开，本文直接使用其合并表，不作任何筛选。

两个预先写死的问题，阈值都是 0.50，事先声明不因结果改动：

- **问题一（越基比例 T）**：人机组合优于“人单干与 AI 单干中较好的那一个”的效应量占比。写死的判负条款是：**若 $T \geq 0.50$ ，则“越过最强单干很少见”这条支柱被削弱，本文照发，并放宽显影窗口的说法。**
- **问题二（跨条件设计占比 W）**：106 项实验中，有多少采用了同一批被试跨不同搭配条件的设计。写死的判负条款是：**若 $W \geq 0.50$ ，则本文“整片评价记录都缺这种形状”的说法在研究场景一侧判负，结论范围必须收窄，本文照发。**若数据里根本没有设计字段，则记“跑不动”，不得反过来说问题二得到了支持。

只跑一遍，不事后改编码，结果无论朝哪个方向都进正文。

结果一：越基比例 $T = 42.4\%-42.7\%$ 。三种口径交叉核算：按原始三方均值并按指标方向判断， $158/370 = 42.7\%$ ；按数据作者自己的方向校准列， $157/370 = 42.4\%$ ；按数据作者自己计算的协同变量取正号计数， $157/370 = 42.4\%$ 。三口径收敛。

这个数低于 0.50，支持了“越过最强单干不是常态”这一条。但它必须削弱本文原来打算写的那句更强的话。田里的数字是 12%，覆盖作物是 2%，人机的数字是 42.7%——**差了一个数量级**。所以“像农田里那样罕见”这个类比不成立，本文收回“罕见”二字，改为“不是常态”。这一处修改是数据逼出来的，不是修辞选择。

结果二（本文自己算出来的那个数）：假盈余带 $\approx 29.5\%$ 。同一批数据里，人机组合优于“两者的平均”的比例是 72.2%（ $267/370$ ，作者校准列为 71.9%）。把两个数放在一起：

- 优于两者平均：**72%**
- 优于较好的那一个：**42.7%**
- **中间这条带：约 29.5%（ $109/370$ ）**

这条带上的每一个效应量，都是这样一种情形：**组合比“随便挑一个好，但比“挑最好的那一个”差**。它看起来像盈余，实际上是搭配把一个较弱成员抬到平均线以上的结果。在没有“较好单干”这个基线的账本里——也就是几乎所有真实的教育与用人账本里——**这一整条带都会被误记成个人的功劳**。它是把搭配增益误算到个人账上的**暴露面**，占比接近三成。

这个数是本文自己算出来的。原数据的作者关心的是协同有没有、在什么条件下有；他们没有把“优于平均”与“优于最优”之间这条带单独命名，也没有把它读成一个归属问题。而这套区分——胜过平均与胜过最强，是两件完全不同的事——正是第五节那门农学手艺用了一百年的语法。本文做的是把这套语法整批地搬到人机数据上算了一遍。

结果三：问题二判负。106 项实验中，含跨条件被试内设计的有 55 项，**51.9%**（其中纯跨条件设计 43 项 = 40.6%，混合设计 12 项）。这个数超过了事先写死的 0.50 阈值。

按写死的条款，本文的第三层主张判负一半，必须收窄。收窄后的说法是：

缺少“此人固定、搭配轮换”这种形状的，是制度性的评价账本——成绩单、招聘档案、绩效系统；研究场景已有一半以上采用了这种形状，且此侧尚待独立取证。

这个不利结果其实比原来那句更有用。它说明：**这种记录形状是可行的、成熟的、已经在被大规模使用的**。做实验的人早就知道，要把人的效应与条件的效应分开，就得让同一个

人经历不同条件——这是方法学常识。所以制度账本里没有这一页，不是技术做不到，是制度选择的结果。这一句比原来那句“整片都没有”更难被反驳，也更容易。

一处如实记录的失误。第一遍计算时，判断“越高越好还是越低越好”的代码写错了：数据里这一列的取值是 Up 与 Down，而我按 high/low 的前缀去匹配，导致所有 Down 行判反，第一遍算出的 T 是 39.7%。发现后修正，并用数据作者自带的方向校准列与协同变量交叉核对，三口径收敛到 42.4%–42.7% 才定数。这一节所有数字都是修正后的。

这次核算的代理关系，必须写明。这份语料里的“搭配”是“人 × AI 系统 × 任务条件”，它是本文所说的“搭配”（合作者集合 × 工具及版本 × 任务类别）的一个收窄代理——它没有人类合作者的轮换。而且这份语料因为入选标准强制三方基线齐全，本身就是一个高度筛选过的样本；真实世界没有这个强制，所以三成的暴露面在真实账本里只会更大，不会更小——但这句话是推断，不是这次核算证明的。

三层主张与它们的可倒性，按从最易倒到最稳排列：

1. **最易倒：**常席差是 AI 时代毕业生核心竞争力的正确答案。
2. **居中：**那张两条轴的表，与“假盈余带”这个分析工具。
3. **最稳（已按判负收窄）：**制度性评价账本里，没有“搭配构成 + 轮换记录”这一页。

引用本文时请分层引用。第一层倒了，第二、三层不受影响。

十 · 与最近的几种说法的分界

本文最容易被人说成“这不就是某某已经讲过的吗”。下面把最像的几种——请进来，并且每一处都给出观察到什么就算谁对。

跨组织流动的双向固定效应模型（Abowd、Kramarz 与 Margolis，1999），以及教师与医生的增值评估。这是离本文最近、也最危险的一位。它做的正是本文说的事：让人在不同单位之间流动，从而把“人的效应”与“单位的效应”分开来估计。它比本文早了二十多年，工具成熟，有大规模行政数据支撑。

分界在识别条件上。这套方法能把人与单位分开，靠的是**不同单位之间没有共同的常驻成员**——如果每一家单位里都坐着同一个人，那么“这个人的效应”和“单位的效应”就完全共线，谁也分不出谁。而 AI 常驻做的正是这件事：**它在每一家单位、每一个团队、每一个项目里，塞进了同一个成员。**本文正是从这套方法的失效条件处长出来的，并且补了它完全不设的两层：账本形状（那一页在制度记录里根本不存在，不是估不准，是没数据）与“计入即毁”的动力学（这个量一旦被公示并挂钩使用，它就失效，见第十五节）。

判决性对照：随着工具在各行业进一步同质化，若这类模型估出的“人的效应”方差占比被系统性压缩，且不同行业估出的人效应之间相关性上升（因为大家共享同一个常驻成员），则本文的判断成立；若人效应的方差构成保持稳定，则这套方法足够用，本文多余。

沙普利值（Shapley，1953）。它给出的正是“遍历所有搭配、按边际贡献分账”的公理化方案，看起来正好是常席差要的东西。分界：沙普利值处理的是**给定一个已知的、可查询的产出函数**之后如何分配；本文处理的是这个函数的记录根本没被建立，以及建立它的动作本身会改变它。**判决对照：**在能查询到全部子集产出的场合（如可重复的算法评测），沙普利值够用，本文的条款不增加任何东西；在账页缺席的场合（几乎所有雇佣关系），沙普利值无从定义，而本文说的正是那里。

团队生产与监督 (Alchian 与 Demsetz, 1972)。他们比谁都早地说清了：团队里各人的边际产出不可分别计量，所以要请一个人来监督，并让他拿剩余索取权。分界：他们的方案是**加强监督**，本文的方案是**改变记录形状**。**判决对照**：若提高监督强度（更细的过程留痕、更密的检查点）能恢复个人可读性，他们对；若必须靠“让人固定、其余轮换”才能恢复，本文对。经验上，过程留痕在混产条件下恰恰读不出作者，这是本文押的一边。

雇主学习模型 (Farber 与 Gibbons ; Altonji 与 Pierret)。它说：市场一开始靠文凭这类信号，随着共事回合增加，会逐步学到人的真实能力，文凭的权重下降。这个模型看起来能吸收本文——多共事几年不就读出来了吗？分界在噪声结构上：它假设每一回合的读数等于“人的信号 + 独立噪声”，所以取平均就能收敛。而本文说，每回合的读数里有一块是**常驻共有成分**，它不随回合独立变化，平均不掉。**判决对照**：让一个人长期待在固定的搭配里（同一批同事、同一套工具、同一类任务），累积再多回合，对他个人评估的方差只会收敛到搭配层，不会收敛到个人；反之若个人评估在固定搭配里也持续变精确，则雇主学习模型够用。

概化理论 (Cronbach 等, 1972)。它把评分方差拆给不同的“面”——评分者、场合、题目——并算出在何种设计下测量足够可靠。它太像本文了：它也在处理“读数不只属于被测者”。分界很硬：概化理论把这些面当作**可抽样的误差宇宙**，目标仍然是把被测者的宇宙分估出来；而常驻同产者不是一个可抽样的面——**它在每一次抽样中都在场，因此与被测者完全共线**。**判决对照**：当搭配不轮换时，概化系数会照常算出一个漂亮的数，不会有任何异常，而 g_r 在此无定义。一个报警，另一个不报警，这就是分界的位置。

明星业绩的不可携带性 (Groysberg, 2010)。他跟踪了华尔街的明星分析师跳槽，发现业绩大幅下滑且多年不恢复——业绩带着团队与平台，带不走。这是本文的近亲，方向也一致。分界：他没有“常驻匿名同产者”这个前提（他的时代还没有），也没有给出可操作的读数协议；他证明了业绩有搭配性，没有给出在承认搭配性之后如何仍然读出人的办法。本文补的正是这一步。

群体智力因子 (Woolley 等, 2010)。他们发现群体存在一个稳定的、跨任务的表现因子，且它与成员个人智力的相关很弱。这条结论对本文极其有利——它说明搭配层有一个真实的、可测的量。但他们测的是**组**，并且不回到人身上；本文要的正是“回到人身上”的那一步，而那一步需要的是同一个人跨多个组的记录，那不是他们的设计。

社交技能溢价 (Deming, 2017)，以及“独立思考”“判断力”这一族实务答案。这是当前最主流的一类回答：既然常规认知任务被自动化，回报就转向难以自动化的社交与协作技能。分界：这仍然是一个**存量式**答案，仍然要靠量表或问卷来测，仍然可以进课程、进清单。**判决对照**：在控制住这类技能的测得分之后， g_r 是否仍能独立预测跨搭配的表现？若不能，存量论够用，本文多余；若能，说明存在一份不被任何存量量表捕捉的、只在轮换中显影的东西。这是本文最愿意接受检验的一条。

信号理论 (Spence, 1973) 与执照/准入制度。教育的价值很大一部分在分选与发信号。分界：本文完全承认这一点，而且第四节就是它的制度版。信号管的是**准入**——在没有共事记录之前，用人方需要一个事前判据。本文管的是**读人**——共事记录开始积累之后，靠什么把这个人和他的搭配分开。两账并行，本文不主张废除凭证（这一点在第十三节还要说，因为它削掉了本文最想要的一个处方）。

遮蔽招聘与工作样本 (Goldin 与 Rouse, 2000 ; Schmidt 与 Hunter, 1998)。幕布后面试音提高了女性入选比例；工作样本测试在预测效度排序里名列前茅。这两项是“遮蔽 + 试做”这条路在人身上的最好成绩。分界：它们借的是遮蔽与试做的形式，没有借感官评价那一套关于**读数归属**的判据装置——没有重复样本核验，没有场次效应分解，

没有描述与偏好的分离。本文借的正是那套装置，并且用它自己的实测得到了相反方向的结论：**遮蔽提高的是程序公平，不解决读数归谁的问题。**一次遮标试音仍然是单回合。

预测机器与判断互补论 (Agrawal、Gans 与 Goldfarb, 2018)。他们说，AI 让预测变便宜，于是预测的互补品——判断——变得更值钱。分界在形状上：他们说的是**内容**（人剩下的活儿是哪一类），本文说的是**读数**（无论人剩下哪一类活儿，你从一次成品里都读不出他）。两者可以同真：即便判断确实是人的活儿，“这个人的判断好不好”仍然只能在跨搭配的差分里读出来，仍然读不出于单件。

利益攸关 (Taleb, 2018)。有下注、承担后果的人才可信。分界：当所有人都以同样方式承担敞口时，这条不能分辨人；而 *g_r* 仍然可以。他管的是激励与筛选，本文管的是读数的可分离性——一个是要不要信，一个是能不能读。

汇聚系列。上面这些，从产权、管理实证、心理测量、计量经济、人机实验五个方向，逼近同一件事：**一次读数、一个东家、一个组，都不回到人。**每一位都差同步：**没有人把基线当作一个常驻的匿名同产者**，因此都还相信“多测几回就能收敛到人”。这一步一旦迈出，“多测几回”就不再是提高精度的问题，而是“测的到底是什么”的问题。

十一 · 与几项相邻判断的分界

还有一批判断，与本文出自同一片问题域，写作时间也相近，必须逐一分清楚——它们最容易被当成同一件事的不同说法。

关于换手率的那一组判断。有一组研究处理的是“位置上的人换了”这件事：一个位置上的承担者多久换一次，换手之后责任如何接续，谁在换手的间隙里落空。它读的是**流动本身**——谁离开了、谁接上了、有没有接上。本文读的是**反面**：以那个不流动的人为参照，对流动的部分做差分。两个读数是正交的：同一份记录里，换手率很高而常席差很低（人来人往，谁在都一样），或者换手率很高而常席差很高（人来人往，只有他在时不一样），都完全可能。请不要把“轮换”这个词的出现当成同一件事——一个把轮换当被观测对象，一个把轮换当观测手段。

关于发起权与首动次序的那一组判断。它们处理“谁先开口”“发起权由谁持有”“收回发起权”这类问题，读的是一个动作序列里的位置先后。本文不涉及次序，它涉及的是同一位置上跨多次搭配的差分。判决对照：一个人可以从不首动而常席差很高（他不发起，但他在的组就是不一样），也可以次次首动而常席差为零。

关于亲历与归属真值的那一组判断。有一条判断处理“这件事到底是不是长在这个人身上”——形成的过程属不属于他，还是他只是挂了名。它与本文的关系最需要说清楚，因为看起来最像。分界是：**即便归属真值全知**——假设有一个全知视角，能确切告诉我们这件作品有几成是他的——**单回合的读数仍然不载这个信息**。也就是说，那一组判断处理的是上游的“发生归属”，本文处理的是下游的“读数可分离性”。两件事互不蕴含：归属可以完全清楚而读数完全不可分离（他确实做了七成，但没有任何记录形状能让评价者读出这七成），也可以归属含混而读数可分离（谁做的说不清，但只要他在，成绩就高一档）。

关于在场与完成资格的那一组判断。有一条判断处理“一件事有没有真的被经历过”“在场是不是可以被代签”。它管的是**存在资格**——这件事算不算真的发生在这个人身上。本文不问这个，本文假定它已经发生，只问它如何被读出与定价。上游的发生 vs 下游的读数，这条线在本节里出现第二次，因为它是本文与整片相邻判断之间最主要的一道分界。

关于认知发生权的那一条判断。它处理的是：在答案变得极其便宜之后，一个人还有没有机会让认知真的在自己身上发生一次——过早拿到答案会取消发生的机会。它管的是**能力**

的发生，本文管的是**能力的读出与定价**。两者可以互相为难：一个人可能认知确实在他身上发生了（发生权保住了），但因为账本形状，谁也读不出来，于是不被定价；反过来，一个人也可能什么都没发生，却因为搭配好而读数漂亮。**这两条判断合起来才是完整的一对**：一条管别让它不发生，一条管发生了要读得出来。

关于评价尺度与自我观众的那一条判断。它处理的是：当评价的标度被取消或变得模糊，人会如何在心里安置一个想象中的观众，并因此改变行为。它读的是被评者的**内部状态**，本文读的是**账本的形状**。判决对照：一个人可以完全没有自我观众而常席差很高，也可以自我观众极强而常席差为零。但两者有一处真实的交接：本文第十四节提出的记账改造，会显著改变被评者心里那个观众的形状——从“我这一件要好看”变成“我在哪几种搭配里都要在”，这个变化的心理后果不在本文能力范围内，应由那一条判断继续处理。

关于课堂评价承重转换的那一条判断。它处理的是评价在 AI 环境下如何从“信息搬运的核验”转向“发生的识别”，以及制度为什么迟迟不动。它与本文的分工是：它管**评价要识别什么**，本文管**账本要长成什么形状才可能识别**。两条判断在“制度延迟”这一点上完全一致，且互为佐证——本文第九节判负后收窄的那句话（研究场景已有此形状，制度账本没有），正是它所说的制度延迟的一个可量化侧面。

十二 · 那块共同的地基

再往下挖一层。上一节请进来的那些说法互相争吵得很厉害，但它们全都站在同一块地上：

一份关于个体的读数，其中非噪声的那一部分，属于被读的那个个体。

真分数理论这样说（观测分 = 真分 + 误差，重复即收敛）；概化理论这样说（把各个面拆开之后，剩下的宇宙分是他的）；雇主学习这样说（回合足够多，市场就学到他的真能力）；增值评估这样说（扣掉环境与起点，剩下的是这个老师/医生的贡献）。它们争的是怎么把噪声剥干净，从没有人怀疑剥干净之后剩下的那一块有主。

承认每一回合的生产里都常驻一个匿名同产者，这块地基就塌了一角：**剥干净噪声之后剩下的那一块，一部分是这个人的，一部分是那个所有人共有的成员的，而后者恰恰不是噪声——它是系统性的、稳定的、可复现的，因此它不会被重复消掉，只会被重复加固。**读得越多、测得越准、一致性越高，它在读数里的份额越稳固。

这就是为什么这件事不能靠“把测量做得更好”来解决。测量做得越好，共有成分被读得越准。要动的不是测量，是**记账的单位**——第四节那门手艺一百年前就动过一次（从一次品评改成一个持有者的跨年份记录），代价是几十年的立法与制度建设。

十三 · 不适用的边界

本文有两处被材料真正削弱的地方，必须写在正文里。

削弱之一，来自田里的数据：不要指望常席差普遍为正，也不要“用”罕见这个词。第九节的核算把“越基罕见”这个类比打掉了一半：农田是 12%，覆盖作物是 2%，人机是 42.7%，差一个数量级。所以人机搭配比农业搭配更容易产生真盈余，这一点必须承认。同时另一头也要说清：**常席差在多数人身上、多数任务上，接近零或为负。**这不是一个人人都有、只是没被测出来的隐藏优点。它是一个分布，多数在零附近。因此本文的适用范围收窄为：**那些搭配读数确实被采用的评价场合——也就是有人要根据混产成果来判断个人的场合。**一个人独自完成、无搭配的工作（还有多少呢），本文不适用。

削弱之二，来自那块地的规程：公示义务不能废，本文最想要的那个处方被拿掉了。顺着本文的逻辑，最自然的处方是"废除单件考核、废除简历筛选、一切改成跨搭配记录"。这个处方不成立，理由是第四节那一家一百年前就给出的：**市场需要事前信号，不能等到 N 个回合之后。**一个用人方面对两千份简历，没有任何跨搭配记录可用，凭证是他唯一能用的东西。凭证管准入，本文管读人，**两账并行，谁也不能取消谁。**而且要老实说：建立跨回合账本的成本极高——原产地规程从立法到形成可撤级的追索机制，用了几十年，中间经历过丑闻、诉讼与分类重建。指望教育与用人系统在五年内长出这一页，是不现实的。这一处削弱动了本文的价值排序：本文不再主张"取消"，只主张"并行增设"。

另外三条边界，说明白省得被误用：

- **搭配不足三个的人，g_r 无定义。**应届生绝大多数落在这一格。所以本文**不能被**用来筛应届生——这是本文最容易被滥用的方向，也是第十五节要专门写的原因。本文的实践含义是：给应届生**创造**跨搭配的记录机会（轮岗、多项目、多队友），而不是拿一个算不出来的数去卡他。
- **基线不可估的场合不适用。**全新的任务类型、没有同类历史的岗位，b_j 估不出来，整个读数无从谈起，不得用行业平均或任何默认值凑。
- **安全关键系统不适用于为读数而制造轮换。**见下一节。

十四 · 能读出人的账本长什么样，以及它为什么多半不会被建

如果要建那一页，它长这样，一共四栏：

1. **搭配构成：**这一次生产里有谁——人的名单，以及工具及其主版本。工具版本必须记，因为它决定基线。
2. **轮换标记：**这个人的上一次记录里，搭配是不是同一个；哪一项变了。
3. **基线估计：**这一类搭配、这一类任务，通常做到什么水平。这一栏最贵，因为要有可比的历史。
4. **逆境标记：**这一回合是不是坏年景——时间压力、信息不足、需求变更、事故处置。第五节那门手艺早就发现，搭配的真实差别在坏年景才放大；顺境里最强单干往往就赢了。

有了这四栏，g_r 可算；没有这四栏，怎么算都算不出来。

现实里最接近这一页的三种记录，各差一点：开源贡献史（缺工具版本栏）、多项目制的咨询与工程组织（有搭配构成，缺基线栏）、以及研究实验（四栏基本齐全，但只在实验里，不在雇佣关系里）。**第九节判负的那个结果在这里发挥了它最大的用处：**既然一半以上的实验早就在用这种设计，那么"技术上做不到"这个借口不成立。剩下的只能是制度选择。

为什么多半不会被建？三个理由，一个比一个硬。

第一，**它读起来慢。**至少三个搭配才有数，坏年景才显影，可能要好几年。而招聘要在两周内做完。

第二，**它触碰隐私与协作关系。**记录"谁和谁一起做的、用的哪一版工具"，等于给每个人建一份协作图谱。这件事的滥用风险极高，第十五节要专门处理。

第三，也是最根本的：**它一旦被建成并公开挂钩，就会失效。**这一条不是风险，是这个量的性质。下一节详述。

十五 · 伦理与证据边界

本文交出的是一个比率，一个可以拿去考核人的数。所以下面三条是硬条款，不是建议。**三条缺一，这个读数就不该被使用。**

第一，这个读数指的是位置，不是人。 g_r 描述的是"某人处在跨搭配不变项这个位置上时，记录呈现出的差分"。它不是一个人格特质，不是一种能力，不是天赋。同一个人换一片领域， g_r 可以完全不同——因为基线变了、搭配类型变了。任何把 g_r 当成"这个人的等级"的用法，都是误用。它更接近一个位置读数，而不是一张证书。

第二，不得为了把这个数做好看，而在安全关键系统里安排轮换或制造变动。这一条针对的是一个非常现实的危险：本文说"要轮换才读得出人"，管理者听成"那就多轮换"。在手术团队、飞行机组、电力调度、核安全这类地方，**固定搭配本身就是安全资产**——磨合出来的默契、共享的心智模型、无需言语的配合，是事故率下降的直接来源。为了让某个人的读数可算而拆散一个磨合好的团队，是拿安全换测量。**在安全关键系统里，本文的读数应当直接判为不适用，而不是"谨慎使用"。**

第三，已经按这个读数做出不利安排的，必须公开编码规则并给出复核通道。如果有机构用 g_r 做了不续约、不晋升、不录用的决定，它必须公开：搭配三元组怎么划的、基线 b_j 怎么估的、最少搭配数是几、哪些回合被判为不适用。并且必须给当事人一条复核通道——尤其是"搭配不足三个而被按零处理"这种情形，那是定义错误，不是低分。

反噬预言，写在正文里。这个读数一旦进入考核，最省事的达标办法是什么？不是提高常席差，而是在**文书里做手脚**：把不轮换的搭配记录成轮换（换个项目代号即可）、挑选自己表现好的回合上报、把坏年景的回合标成顺境不计入。更深一层：一旦 g_r 被折算成一个可携带、可公示的分数，两件事同时发生——它成为一个公开的标签，于是第三节那条"期望改写读数"的机制重新启动，评价者看到高 g_r 的人，读到的东西真的不一样了；同时它成为那个匿名常驻同产者的又一个训练目标，被优化、被均摊、被追平。**这个数在被广泛采用的那一刻，就失去了它本来要携带的关于人的信息。**

这与"一个指标一旦成为目标就不再是好指标"那条常被引用的规律**不是同一回事**，分界可以判：那条规律的机制是激励——因为有奖惩，所以人去优化指标。本文这条的机制不需要激励：**即便完全不挂钩任何奖惩，只是把这个数登记并公示，它照样衰减。**判决形式很干净——设两组，一组只登记不公示不挂钩，一组公示但不挂钩：若公示组的分数与后续实际表现的相关性衰减快于只登记组，则本文这条成立（污染机制）；若只有在挂钩奖惩时才衰减，则那条老规律够用，本文这条多余。

我要老实说：**我看不到出口。**任何让常席差被广泛用起来的方案，都要求它被记录、被计算、被传达；而记录与传达本身就在毁它。可能的缓解只有三种，都不彻底：让它只在机构内部使用不对外公示；让它只作为"值得多看两眼"的触发器而不作为排序依据；以及让基线 b_j 保持不公开，使优化没有靶子。三种都靠制度自律，都会被绕过。这一条是本文最弱的地方，我不掩饰。

十六 · 什么情况下本文算错

按第七节那一册口径，本文的可判负条件如下，每一条都写明观察什么、以及什么不算数。

判负条件一：在控制住既有的能力量表（认知能力测验、社交能量量表、结构化面试评分）之后， g_r 对跨搭配表现没有独立的预测力。观察对象：任何拥有三个以上搭配记录的人群。若成立，则存量式答案够用，本文的第一层主张倒。

判负条件二：在固定搭配里累积足够多的回合之后，对个人的评估精度持续提高，而不是收敛到搭配层。若成立，则雇主学习模型够用，"共有成分不会被平均掉"这条机制不成立。

判负条件三：跨组织流动的双向固定效应模型在工具高度同质化之后，估出的人效应方差占比保持稳定、跨行业人效应相关性不上升。若成立，本文说的那种共线并没有发生。

判负条件四：把 g_r 只登记不公示、不挂钩奖惩，它与后续实际表现的相关性仍与公示组一样衰减——那么第十五节那条"不需要激励也会衰减"的机制不成立。

判负条件五（已部分实现）：第九节问题二已经判负，第三层主张已按条款收窄。若进一步的证据表明，制度性评价账本里其实已有相当比例记录了搭配构成与工具版本，则第三层主张完全倒。

写死日期的那一条，见第十八节。

十七 · 本文自己

用本文自己的判据审本文自己，不是为了谦虚，是因为本文如果不能给自己定位，它就不该被别人用来定位任何人。

本文是一次单回合的混产。它由一个人与一套模型共同完成，模型的具体版本、参与的深度、每一节各改了几遍，都没有留下记录。按本文自己的判断，读者从这一件成品里读到的，是"这位作者 × 这套工具 × 这一类题目"这个搭配，读不出其中任何一方。**本文正坐在自己描述的靶格里。**

这不是一句自谦，它有具体后果，三条：

后果一：本文的任何单次评价都不可信，包括好评。如果有人读完说"这篇很有见地"，按本文自己的主张，这个判断读的是搭配，不是作者。所以本文不能拿任何一次好评当作它主张成立的证据——那是自相矛盾。

后果二：本文唯一能指认的自身可读格样本，是同一执行者跨搭配的记录史。这位作者在同一套工作方式下，跨越几十个不同题域、不同工具版本、不同合作形态，产出过一批可比的东西。**那批记录才是本文自己那一页账。**如果那批记录里，成绩的差分并不稳定偏向这位作者——也就是说换个人用同一套方法能做出同样水平的东西——那么本文关于自己的那部分主张就倒了。这一条我无法自证，只能指出证据在哪儿、该怎么查。

后果三：本文提出的这个读数，本文自己没有资格用它给任何人打分。参与写作的一方不能为自己参与的文本发认证分，这是一条一般原则，在这里有加倍的理由：本文正是在论证"参与者读不出参与者"。任何关于本文水准的判定，都必须由独立的、未参与写作的一方给出，并且最好不止一次、不止一组。

还有第四条，关于本文的处境：**本文所主张的那种可读性，本文自己享受不到。**因为一篇文章就是一件单件作品，它被引用、被评价、被排序的方式，全部是单回合的。本文要求的记账形状，正是本文这个文体所不具备的。这一点不构成对本文内容的反驳，但它确实说明：**改变账本形状这件事，靠写文章推不动**——写文章本身就在旧账本里进行。

十八 · 结论，与一个写死日期的赌注

回到那三张成绩单。

第一张读出的是搭配。第三张诚实地承认自己读的是一个组。而第二张纸上那件真正值钱的东西——三年里，无论跟谁组队、做什么题、用哪一版工具，那个组交出来的东西总是好一点点——**成绩单上没有为它设一栏**，所以它不被读出、不被定价、不被承认。

AI 时代大学毕业生的核心竞争力，就是那件没有栏目的东西：**常席差**。它不是一件作品的品质，不是一纸出身的凭证，也不是一条能写进培养方案的能力条目。它是“别的都换过之后，还剩下的那道差”。**在单回合的评价记录里，它不是难测，是没有定义。**

所以三句实践含义：

对学生：不要把力气全花在把单件做得更漂亮上——那一格的分辨率正在归零。去积累一件事：**在尽可能多的不同搭配里留下可比的记录**。不同的队友、不同的任务、不同的工具版本，最好还有几次坏年景。这不是为了简历更花哨，是为了让你身上那道差有机会被算出来。

对学校：与其把资源继续投进监考、检测与更严的单件闸门（那是靶格），不如做一件更笨的事——**在培养记录里加上四栏：搭配构成、轮换标记、基线、逆境标记**。这四栏一旦有了，很多现在吵不清的事都会自己清楚。

对用人方：凭证继续用，它管准入，本文不反对。但请把“三个以上不同搭配下的可比记录”当成一个独立的、比作品集更值钱的东西去索取和保存。也请记住第十五节那三条，尤其是不要为了让这个数好看去拆散安全关键的固定团队。

最后，一个写死日期的赌注。

到 2031 年 12 月 31 日为止，主流的招聘系统、绩效系统与学籍成绩系统，不会新增“搭配构成（协作者 + 工具及其主版本）与轮换记录”这一类结构化字段。

什么算我输：任一主流人力资源系统或学籍系统，在其标准字段里增加了“本次产出的协作者名单”与“所用工具及主版本”，且这两栏被用于个人层面的评估或定价。

什么不算我输——这一条必须写死，否则赌注是空的：**仅仅新增“是否使用了 AI”这样一个二值披露字段，不算命中**。这类字段现在已经在出现，它记录的是“用没用”，不是“和谁、用哪一版、在什么基线下”。二值披露解决的是诚信问题，不是可读性问题；它甚至会加固靶格，因为它把注意力引向“这东西干不干净”，而不是“这个人读不读得出”。

如果 2031 年底这一栏真的出现了，本文的第三层主张倒，而我很高兴地认输——因为那意味着第二张成绩单上的那个孩子，终于有了一栏可以待的地方。

注释与说明

关于第九节所用数据。数据为公开发表的人机协作系统综述与元分析的开放数据（106 项实验，370 个效应量），由后续再分析者整理并公开于开放科学框架平台。本文使用其合并数据表，未作任何样本筛选。计算规则（两个问题、0.50 阈值、判负条款）在取数之前写定，仅执行一遍，未因结果调整编码。方向判定的一处编码失误已在正文第九节写明，并已用数据作者自带的方向校准列与协同变量交叉复核。

关于三家材料的引用性质。第三、四、五节所引的三门手艺的核心判断句，均取自各自领域的原始文献或法规条文，不是本文代拟。三处最要紧的实证——重复样本的评委一致性、标签对体验层面的改写、混作越过最强单作的比例与显影年限——均可在所引文献中逐条核对。

字数与版本：正文约 2.0 万汉字。版本 1.0，2026 年 9 月 2 日。

人机分工声明：本文由作者与一套大型语言模型共同完成。选题、判断的取舍、三家材料的选定、可判负条款的写定与最终定稿由作者负责；文献检索、数据下载与计算、初稿撰写由模型执行。第九节的计算脚本与原始输出可备查。按本文第十七节的判断，读者从本文这一件成品中读到的是这一搭配的产出，而非其中任一方；本文不为自己主张任何评价等级，任何关于本文水准的判定应由未参与写作的一方独立作出。

参考文献

Abowd, J. M., Kramarz, F., & Margolis, D. N. (1999). High wage workers and high wage firms. *Econometrica*, 67(2), 251–333.

Agrawal, A., Gans, J., & Goldfarb, A. (2018). *Prediction Machines: The Simple Economics of Artificial Intelligence*. Harvard Business Review Press.

Alchian, A. A., & Demsetz, H. (1972). Production, information costs, and economic organization. *American Economic Review*, 62(5), 777–795.

Altonji, J. G., & Pierret, C. R. (2001). Employer learning and statistical discrimination. *Quarterly Journal of Economics*, 116(1), 313–350.

Amerine, M. A., Pangborn, R. M., & Roessler, E. B. (1965). *Principles of Sensory Evaluation of Food*. Academic Press.

Cardinale, B. J., et al. (2007). Impacts of plant diversity on biomass production increase through time because of species complementarity. *PNAS*, 104(46), 18123–18128.

Cronbach, L. J., Gleser, G. C., Nanda, H., & Rajaratnam, N. (1972). *The Dependability of Behavioral Measurements*. Wiley.

Deming, D. J. (2017). The growing importance of social skills in the labor market. *Quarterly Journal of Economics*, 132(4), 1593–1640.

Dion, R. (1959). *Histoire de la vigne et du vin en France des origines au XIXe siècle*. Paris.

European Union (2012). Regulation (EU) No 1151/2012 on quality schemes for agricultural products and foodstuffs, Article 5.

Farber, H. S., & Gibbons, R. (1996). Learning and wage dynamics. *Quarterly Journal of Economics*, 111(4), 1007–1047.

Florence, A. M., & McGuire, A. M. (2020). Do diverse cover crop mixtures perform better than monocultures? A systematic review. *Agronomy Journal*, 112(5), 3513–3534.

Goldin, C., & Rouse, C. (2000). Orchestrating impartiality: The impact of "blind" auditions on female musicians. *American Economic Review*, 90(4), 715–741.

Groysberg, B. (2010). *Chasing Stars: The Myth of Talent and the Portability of Performance*. Princeton University Press.

Hodgson, R. T. (2008). An examination of judge reliability at a major U.S. wine competition. *Journal of Wine Economics*, 3(2), 105–113.

Hodgson, R. T. (2009). An analysis of the concordance among 13 U.S. wine competitions. *Journal of Wine Economics*, 4(1), 1–9.

Lawless, H. T., & Heymann, H. (2010). *Sensory Evaluation of Food: Principles and Practices* (2nd ed.). Springer.

Loreau, M., & Hector, A. (2001). Partitioning selection and complementarity in biodiversity experiments. *Nature*, 412, 72–76.

Matthews, M. A. (2015). *Terroir and Other Myths of Winegrowing*. University of California Press.

Plassmann, H., O'Doherty, J., Shiv, B., & Rangel, A. (2008). Marketing actions can modulate neural representations of experienced pleasantness. *PNAS*, 105(3), 1050–1054.

Schmidt, F. L., & Hunter, J. E. (1998). The validity and utility of selection methods in personnel psychology. *Psychological Bulletin*, 124(2), 262–274.

Shapley, L. S. (1953). A value for n -person games. In Kuhn & Tucker (eds.), *Contributions to the Theory of Games II*. Princeton University Press.

Spence, M. (1973). Job market signaling. *Quarterly Journal of Economics*, 87(3), 355–374.

Taleb, N. N. (2018). *Skin in the Game: Hidden Asymmetries in Daily Life*. Random House.

Trenbath, B. R. (1974). Biomass productivity of mixtures. *Advances in Agronomy*, 26, 177–210.

Vaccaro, M., Almaatouq, A., & Malone, T. (2024). When combinations of humans and AI are useful: A systematic review and meta-analysis. *Nature Human Behaviour*, 8, 2293–2303.

Vandermeer, J. (1989). *The Ecology of Intercropping*. Cambridge University Press.

de Wit, C. T. (1960). On Competition. *Verslagen van Landbouwkundige Onderzoekingen* 66.8.

Woolley, A. W., Chabris, C. F., Pentland, A., Hashmi, N., & Malone, T. W. (2010). Evidence for a collective intelligence factor in the performance of human groups. *Science*, 330, 686–688.

第十一章 峰兑

【章首】本章给可读条件立第二条：**事件处**。当 AI 把一切可展示产出的地板抬到同一高度，候选人之间的真实差距缩进量具的判等差额之内——量具在灵敏度意义上死亡，而它照常出分，于是死得无声。人身价值里那份平时读数恒为零、只能被稀发事件整笔结算的，本章命名为峰兑；它在产出计费型劳动市场被结构性欠付，欠付、减产与账面改善三件事由劳动侧前提四步推出，电力经济学的“缺钱问题”只是同一结局在另一颗行星上的既有落地。读数是盲读率 m 。本章的实算给了可读条件一个触目的数：同一张分数表上，高分段每分不足十人、中段每分近两千人——中段一分裁定约一千九百个并列者。本章的三条预注册预测全部命中，三处不利侧写照登。本章欠的账：与《底流》共用的那场峰时事件双组实验，未跑；兑付记录字段规范已拟至第零点一版，无一处真实落地。

大白话：平时看不出的那一份，出事那天一次结清

一家餐馆，平时的家常菜，客人吃不出哪个厨师做的。两个厨师，谁在灶上，评分都差不多。

直到某天晚上出事：五十人的宴席临时加了两桌、灶台起火、一位客人鱼刺卡喉。那一晚过去之后，两个厨师的差别，一次性全部显形。

《峰兑》讲的就是这一份价值：**它平时的读数恒为零，只在稀发事件处整笔结算。**

为什么“平时读数为零”不等于“它不存在”？

这是这一章最要紧、也最容易被误解的一点。平时读数为零，不是因为量具不够精，而是**这份价值的性质就是如此**——它只在事件处兑付，平时没有可读的地方。你用平时的量具去读它，读到零是必然的，读一万次也是零。

于是就有了一个残酷的推论：**在一个只用平时量具做分配的系统里，这一份价值被结构性地欠付。**平时读得越勤，欠得越稳。

量具是怎么死的？

书里有一个说法很硬：当 AI 把所有人可展示的产出都抬到同一地板，候选人之间的真实差距缩进了量具的“判等差额”之内——**量具在灵敏度意义上死亡了，而它照常出分。**

这就是它可怕的地方：一把已经分不出人的尺子，不会报错，不会停机，它每天照常给出精确到小数点后一位的分数。所有人都以为在测量，其实在读噪音。

书里给了一个触目的实算：在一张常见的分数表上，高分段每一分不到十个人，而中段每一分接近两千人——**中段的一分，实际上在裁定约一千九百个并列者的顺序。**这不是测量，是抽签，只是抽签的结果被印成了分数。

读数是什么？

叫**盲读率 m** ，衡量的是平时读不出、事件处才结算的那一份有多大。它的操作化依赖一样东西——**兑付记录**：一条合格的兑付记录要满足四条，即真实对手方、真实赌注、无重来按钮、对手方签收。

这四条里最容易缺的是最后一条。很多人扛过事，但没有任何可查证的痕迹留下来，于是那次兑付在制度上等于没发生过。

一个小场景。

两个客服，平时的工单评分都在 4.8 分上下，看不出差别。某天系统大面积故障，投诉量涨了二十倍。甲稳住了三十几个最难的客户，其中两个后来续了约；乙按流程回复，客户跑了一半。

一个月后绩效评估，两人的月度平均分几乎一样——因为那三天的极端工单被系统标记为“异常数据”剔除了。

这就是欠付的具体样子：不是有人恶意抹杀，是量具的正常运作把最有信息量的那三天当成了噪音清理掉。

它不是什么？

它绝对不是“鼓励制造事件”。这一点书里写成了硬性伦理条款：**不得制造、放任或延迟事件**——纵火的消防员是这类读数最经典的反噬。峰兑只能等，不能造。

再看两个身边的例子。

小区里的水电工。平时你几乎想不起他，一个月也见不到一次。直到某个半夜水管爆了——那一晚他值多少钱，和他平时的“绩效”完全不是一个量级。如果这个小区的物业按“平时的工单数量”给他发钱，那么他这份价值一年到头都被欠着，直到某天他不干了，大家才在下次爆管时结清这笔账。

保险。你交了十年保费什么也没发生，账面上纯亏。这一份价值不是不存在，它只是**只在事件处兑付**。人的承接力和保险的结构完全一样，区别只在于：保险这件事社会已经想明白了，人的这一份还没有——所以保险有精算，而承接力没有账。

三种常见的误读。

第一种：以为“平时读不出”意味着可以不管。正相反：正因为平时读不出，它才需要被记账而不是被测量。所以这一章给的处方不是造一把更精的尺子，是留一本兑付记录。

第二种：以为可以用演练代替。演练能训练动作，但它有重来按钮、没有真实赌注、时刻是安排好的——三条都不满足。演练出色的人在真实事件中崩掉，或者演练平平的人在真实事件中稳住，都不奇怪。

第三种：把它当成鼓励逞英雄。这是最危险的误读，书里把它写成了硬条款：**不得制造、放任或延迟事件**。一个为了积累兑付记录而让风险发生的人，做的是纵火的消防员那件事。峰兑只能等，等不到是常态；等不到的时候，那一份价值仍然在，只是没有到结算窗口。

给自己做个小检查。过去三年，你手上有几条满足这四条的记录：真实对手方、真实赌注、没有重来按钮、对方留下了可查证的确认？多数人的答案是零到一条——不是因为没扛过事，是因为扛完没有留痕。

它和旁边几个概念的区别。

和《底流》的那对相反预测，是全书最便宜的一次裁决，上面已经讲过。

和《零判份》的分工在于两种“读不出”：未见率量的是**世界拿不拿得出**（供给侧），盲读率量的是**量具分不分得开**（测量侧）。两者会同时出问题——世界拿不出同类物、而量具照样分不开人，那正是今天最典型的处境。

和《他手无事账》的关系很有意思：峰时事件天然是他择的（没有人挑事件发生的时间），所以每一次真实事件同时给出两章的读数。反过来不成立：无预告的抽样可以很频繁，但那不构成事件——没有真实赌注，没有整笔结算。

和《单遍段》的区别在稀发与否，前面讲过：不可撤销步天天有，峰时事件很少有。

一句话记住它：这一章管的是**什么时候才结算**——以及那之前的所有读数为什么必然是零。

怎么长出来：方法、技术与流程

个人层面。

第一，**留兑付记录**。你扛过的事，当时就留一句可查证的痕迹：一封邮件、一条工单备注、一次对方的确认。不必邀功，但要留痕——没有对手方签收，那次兑付在账上不存在。

第二，别为了兑付去制造事件，但也**别躲开真正的事件**。事件是这一份价值唯一的结算窗口，躲开等于永远不结算。

大学发生学教育：给学生一个"事件账"。

学校是一个刻意消除事件的地方——课程有大纲、进度可预期、意外被提前排除。这对学习是好事，但它意味着学生四年里可能没有一次真正的峰时经验，也没有任何一条兑付记录。

- **事件账（可携带）**。给每个学生一本简单的记录：凡他在真实的紧急、超载、突发状况中承担了什么，由**对手方**（老师、合作单位、社团、志愿服务的受助方）签一句话确认。四年下来，这本账比任何证书都稀缺，因为几乎没有人有。
- **真实高压的合法来源**。不制造事件，但可以让学生**进入本来就会发生事件的地方**：急救志愿、赛事保障、迎新与大型活动的现场执行、开源项目的线上故障、真实客户的截止日。这些地方的事件不是为教学而造的，是本来就有的。
- **异常数据不得剔除**。这一条是给教务和课程评价系统的：**不要把极端情况下的表现当作异常值清理掉**——那恰恰是唯一有信息量的样本。评价体系应当单独保留一栏"高压情境下的表现"，不参与平均。
- **明确禁令**。培养方案里要写死：不得为获得读数而制造风险，不得在安全关键的实践场景（临床、实验室、工程现场）中撤走冗余以"创造机会"。

教师的动作。

在真实的突发状况发生之后（设备坏了、场地临时变更、合作方鸽了），**不要急着替学生把事情摆平**。这是极少数不需要设计、天然出现的峰时窗口。教师能做的最有价值的事，是在事后给一句确认："这件事你扛住了，我记一笔。"

最容易做坏的地方。

把事件账挂到评奖上。一旦挂钩，就会出现"制造事件"和"抢占事件署名"两种立刻可见的败坏。事件账的正确用途是**归属**——让扛过的事记在扛事的人名下，而不是让它变成一种可竞争的资源。

原文

峰兑——AI时代大学毕业生的核心竞争力，是一份平时读数为零、只在稀发事件处整笔结算的价值；而现行评价与薪酬制度对它存在结构性欠付

摘要 AI 把几乎所有可展示产出的地板抬到同一高度之后，"毕业生的核心竞争力是什么"这道题出现了一个此前不必面对的怪相：所有候选人在所有量具上都读作"差不多"。本文的判断是，这个"差不多"首先是量具的读数，不是人群的读数——当候选人之间的真实差距缩进量具的判等差额之内，量具就在灵敏度意义上死亡，而它照常出分，于是死得无声。据此提出承重命题（三重否定式）：AI 时代毕业生的核心竞争力，不是任何平时量具读得出的显露优势，不是训练路径时长可换算出的凭证价值，也不是平均负载下的产出贡献；而是**峰兑**——一个人价值中平时读数恒为零、只能被稀发事件整笔结算的那一份。峰兑在产出计费型劳动市场中被结构性欠付——欠付、减产与账面改善三件事可由劳动侧四步独立推出，电力经济学的「缺钱问题」只是同一结局在别处被独立发现、并已给出制度解的旁证。材料取自三处彼此无引用往来的实务体系：临床试验的非劣效性判读（含灵敏度与渐端的自查文献）、民航飞行员资质的时数之争、电力容量市场的设计文献。零情态词判据一句：在正式记录里，一个人扛住那次事故所动用的东西，被记在谁名下、算多少？检验按事先注册的三条判负条款执行：成绩量具的顶格份额越线（ $43\% \geq 40\%$ ）、雇主弃用该量具（ $73.3\% \rightarrow 42.1\%$ ）、且越线在先弃用在后（ ≤ 2008 早于 2019 ），三条均获支持，三处不利侧写照登。结论一句：谁先为自己挣到第一张可核验的兑付记录，谁先离开那片所有人都读作一样的盲区。

关键词 峰兑；盲读率；判等差额；量具灵敏度死亡；容量欠付；毕业生就业

English Title: Peak-Settled Worth: Why a Graduate's Core Competitiveness in the AI Era Is the Share of Value That No Within-Window Assay Can Read

Abstract: When AI raises the floor of every displayable output, all graduates read as "about the same" on every instrument. We argue this sameness is a property of the assay, not of the population: differences have shrunk inside the instrument's margin of equivalence, so the assay has died in the sense of ICH E10's *assay sensitivity* while continuing to emit scores. We name the residual object **peak-settled worth**: the share of a person's value that yields zero in every ordinary reading and settles in full only at rare events. We derive from labor-side premises alone that output-metered labor markets structurally underpay this share — converging with the *missing money* theorem of capacity-market economics (Boiteux 1949; Cramton, Ockenfels & Stoft) as an independently discovered isomorph, hence under-produce it, and that the under-production registers as improvement on every periodic metric. We engage and separate from Weisbrod's (1964) option value, Spence's (1973) pooling equilibria, Cyert & March's (1963) organizational slack, Schmidt & Hunter's (1998) selection-validity tradition, Janz's (1982) behavioral interviewing, Filippas, Horton & Golden's reputation inflation, supervisory stress testing, Perrow's *Normal Accidents*, and Taleb's *Antifragile*. A preregistered order test (grade-scale saturation preceding employer abandonment of GPA screening) is reported with all adverse findings.

Keywords: peak-settled worth; margin of equivalence; assay sensitivity; missing money; graduate employability

一 两份一模一样的简历，和一台没人用过的灭火器

先摆两个日常画面。

第一个画面在招聘会上。2026 年的面试官桌上放着二十份应届生简历，每一份都干净、完整、漂亮：项目经历三段、作品链接可点、自我陈述通顺得挑不出一个错字。十年前，这二十份里会有三四份明显好、五六份明显差，剩下的居中；现在，二十份读起来全都在同一个水平线上。面试官的第一反应通常是感叹“现在的学生真强”，第二反应才是那个更冷的疑问：**如果所有人读起来都一样，我手里这套读法还在读出东西吗？**

第二个画面在任何一户人家的厨房角落：一台灭火器。它买来三年，一次没用过，压力表指针一直停在绿区。问它“这三年产出了什么”，答案是零；按任何一种产出台账，它是这间厨房里最没用的东西。可是没有人因此把它扔掉——因为大家都隐约知道，它的全部价值不摊在这三年的每一天里，而是整笔压在某个还没来的下午。**它平时的读数是零，不是因为它的没有价值，而是因为它的价值不按“平时”结算。**

本文要说的是：这两个画面是同一件事的两面。AI 时代毕业生身上真正剩下的那份竞争力，正在变成灭火器式的东西——平时任何量具都读不出来，只在稀发事件处整笔兑付。而我们的招聘、考核、薪酬制度，全部是按“平时”计费的。这中间的缺口不是修辞，它有名字、有机制、有可以数的读数，也有一段已经发生的历史可以查证。

二 三种现成的回答，与它们共同站着的那块地

关于“看起来都一样时怎么分高下”，至少有三个彼此从不引用的实务体系各自给出过成熟答案。

第一个体系来自临床试验。当新药在伦理上不能与安慰剂对照时，就与老药比“不劣于”：只要新药的疗效差距落在事先划定的判等差额（非劣效性界值）之内，就判合格。这套方法学有自己的奠基层与自查传统——监管指南 ICH E10 专门为它立了一个概念叫 **灵敏度**（assay sensitivity）：一场试验必须有能力和探测到既定大小的差距；一场没有灵敏度的试验，会在两样真有差别的东西之间读出“等同”，而且读得理直气壮。自查文献还命名了它的慢性病：**渐蠕**（biocreep）——每一代都只与上一代比“不劣于”，判等差额一代代吞掉一点点真实退步，标准整体缓慢下滑，而每一场单独看都合格。

第二个体系来自民航驾驶舱。2009 年科尔根航空 3407 号班机坠毁后，美国立法把副驾驶的准入门槛提高到 1500 飞行小时——赌注押在“走过的路径本身就是资质”这一边。而国际民航界的另一派推行基于胜任力的训练与多人制机组驾驶员执照（MPL），二百多个小时加模拟机就可上座——赌注押在“展示出来的能力才是资质，路径时长只是代理”。两派吵了十几年吵不出胜负，原因很诚实：真正要防的那类紧急情况太稀发，任何一场考核窗口里都抽不到样。2009 年把 1549 号班机降在哈德逊河上的萨伦伯格机长事后说，他一生都在向经验的银行里存小额存款，1 月 15 日那天是一次巨额支取。**那个“银行账户”平时没有任何仪表读得出余额，它只在支取那一刻显形。**

第三个体系来自电网。电力市场发现过一个著名的结构性缺口：如果只按发出来的电量付钱（产出计费），那么专为极端峰荷备着、常年不转的备用机组就挣不回成本——学界称之为**缺钱问题**（missing money）。峰荷定价的奠基工作可回溯到 Boiteux（1949），而现代容量市场的设计文献（Cramton、Ockenfels 与 Stoft 等）给出的解法是发明**第二种付费**：不为电量付钱，专为“站在那里、随时可发”这件事本身付钱，用拍卖给“待命”定价。

2021 年 2 月的得克萨斯大停电，教科书式地演示了当一个系统长期只按产出计费时，被欠付的备用容量会怎样在最需要它的那个星期集体缺席。

三个体系，三种口音，但它们站在同一块从没被说出口的地上。把它们各自的争论压成一句：它们争的都是“一个载体的价值该由哪一段记录来裁归、裁多少”；而它们共同假定的是——凡真实存在的价值，平时的量具迟早读得出来，读数与价值之间的缺口，总可以由一个事先定好的换算率补上（差额换算、小时换算、容量价换算）。这句话念给三个体系听，它们都会说“这还用说吗”。正因为如此，它值得被推翻一次。

三 来自差额判读自己的观察

推翻它的材料不必外求，就在第一个体系自己的档案里。本文在此处有一个必须写明的双重身份：非劣效性方法学既是本文的供料方，又是被推翻方；用来推翻它的不是任何人的立场，而是它自己发表的观察结果。

替加环素（tigecycline）是一支 2005 年经非劣效性试验获批的广谱抗生素。它一场一场地通过了差额判读——每一场试验里，它与对照药的疗效差距都落在判等差额之内，逐场合格。2012 年，Prasad 等人把 13 场随机试验汇总，读出了单场永远读不出的东西：替加环素组的全因死亡率系统性高于对照组，校正风险差 0.6%，约每治疗 143 人多死 1 人。FDA 于 2010 年发出警示、2013 年加上黑框警告——监管体系里最重的一级。

请注意这件事的形状，它比“一个药出了问题”深一层：每一次期内读数都说“不差”，而差距一直都在；它不在任何一场试验的窗口里显形，只在稀发事件（死亡）跨窗口汇总的地方整笔显形。差额判读体系自己用最贵的方式证明了：读数与价值之间的缺口，恰恰在最要命的那一段上，换算率补不上。合格是逐场的，死亡是累计的。

这个观察一旦拿到手，前面那块共有的地就塌了。塌掉之后站起来的东西，就是本文的承重命题。

四 承重命题：峰兑

AI 时代大学毕业生的核心竞争力——

不是任何平时量具读得出的显露优势。不是因为显露不重要，而是因为 AI 把可展示产出的地板抬进了判等差额：当二十份简历的真实差距全部缩进量具分不出来的那个区间，量具照常出分、照常排序，但排序里已经没有信息。这不是候选人变得相同，是量具在灵敏度意义上死亡——而灵敏度死亡不报错，它安静地继续营业。

不是训练路径时长可换算出的凭证价值。四年、学分、绩点，本质上都是民航的“1500 小时”：拿路径长度代理不可考的储备。这个代理在 AI 之前是粗糙但可用的；AI 介入后，同样的四年里有多少段路是本人走的、有多少段是工具替走的，路径台账上没有字段，小时换算率失锚。

也不是平均负载下的产出贡献。平时的产出正是 AI 溢价最快摊平的地方；一个只按平时产出计费的市场，会把人身上灭火器式的那部分价值计成零——这不是雇主刻薄，是产出计费的算法性质。

而是峰兑：一个人价值中平时读数恒为零、无法被任何期内比较读出、只能被稀发事件整笔结算的那一份。它有四条构成性质：

1. 期内盲：在事件到来之前，任何常规量具（考试、绩效、作品比对、面试）对它的读数与零不可区分；

2. **事件结算**：它的兑付单位是"一次真实事故/一次真实峰荷/一次没有重来按钮的关口"，兑付是整笔的，不可按日摊薄；
3. **不可点播**：不能为了考核它而按需制造一次事件——制造出来的事件（演练、模拟）能考出程序，考不出真实损失承受，因为"没有重来按钮"本身是它的构成条件之一；
4. **欠付即减产**：在产出计费的市场里它挣不到平时工资，于是被理性地减配——而减配在每个季度的常规读数上都表现为成本改善。

第四条是本文最想钉住的一句，第九节将不借任何外域词汇把它推出来：**对峰兑的欠付不是道德失误，是产出计费制度的数学后果；而它造成的短缺，在短缺酿成事件之前，账面上全是利好。**

五 辨别格

把两根轴立起来：横轴是**平时可读性**（现行量具在期内能否读出这份价值），纵轴是**峰时兑付额**（稀发事件到来时它结算多少）。

	峰时兑付额高	峰时兑付额低
平时读得出	① 双显件：平时亮眼、关键时也顶用（既能日常出活又扛过事的老手）	② 演示件（靶格）：一切量具上都漂亮，事件一来无可支取
平时读不出	③ 峰兑件（本文对象）：期内读数为零，事件处整笔兑付	④ 空件：平时没有读数，峰时也没有东西

靶格是②，演示件。它必须能通过全部形式审查，否则不配当靶格——而它确实能：简历完美、笔试满分、作品链接全部可点。它的三条签名逐条对上：

- **短期指标表现为改善**：一个组织把招聘与考核对齐到量具可读的那一列，每一季的达标率、人均产出、录用者画像都在变好；
- **失效不产生信号**：演示件的空账在事件到来之前没有任何仪表会响——正如替加环素在汇总之前场场合格；
- **自我加固**：量具越密（更多测评、更细 KPI），对齐越彻底，②格占比越高；正规化不是解药，是病程。

②与③在一切平时记录上不可区分——这正是问题所在，也是本文全部读数要去测的东西。

六 读数：盲读率，与那个还不存在的字段

主读数：盲读率 m 。

- 定义：在一个筛选场次中，用现行主量具对候选人做两两比较，差距落在该量具判等差额之内（即无法据以定序）的比较对，占全部比较对的比例。 N 的取法：同一职位进入同一轮评审的候选人集合的 $C(n,2)$ 对。
- 判等差额的取法（写死一种，不许场内换）：该量具在前 AI 基线期（本文取 2015–2019）能稳定复现排序的最小分差；无基线记录的量具，取评分者复评标准差的 2 倍。
- 粒度与边界例：两位候选人总分差恰好等于差额，计为“分得开”（不入盲读对）；同分并列计入盲读对。一次“比较”以主量具总分为准，不合并多个量具（合并规则另立即另立一份手册）。
- 编码规则：① 只取进入同轮的候选人，不跨轮拼对；② 缺分者剔除，剔除数须报告；③ 复评分差取同一评分者间隔两周的复评；④ 差额一经取定，全篇（含证伪节与赌注节）引用同一数值。
- 空表头：场次编号 | 职位 | n | $C(n,2)$ | 差额 | 盲读对数 | m 。
- 不适用： $n \leq 4$ 的场次；主量具为纯排名制（强制区分）的场次——排名制强行产生序，但序里是否有信息须另测（复评排序相关），不在 m 的口径内。
- 三处引用一致性自查：正文本节 = 第十五节证伪条款 F2 = 第十五节赌注，同一定义同一差额。

这套口径不是纸上的。拿一把真实量具做一次演示性实算（演示，非假设检验，不进证伪节）：2025 年河南高考物理类一分一段表（省教育考试院发布）给出每个分数上的并列人数与累计位次。按官方锚点算每分密度：700–710 分段约 9.6 人/分，666–682 分段约 152 人/分，600–650 分段约 752 人/分，而 460–500 分段约 **1,946 人/分**——同一把尺子，顶端与密集区的分辨率相差两百倍。在四十六万人聚集的中段，量具最细的刻度（1 分）每出一次分，就在约一千九百个它分不开的人之间画了一条录取边界；一个 5 分窗里站着近万人。官方表格自己写着分数相同者为并列，随后的录取规则再用单科排序把并列强行拆开——这正是第八节那句“有序不等于有信息”在国家级量具上的样子。AI 抬升地板对毕业生量具做的事，就是把每一把尺子的读数分布推向这张表的中段。

本文对 m 的预言是可裁决的：随 AI 渗透率上升，同类场次的 m 单调上升；且雇主端将观察到“量具更替加速”——一种笔试或作业形式从采用到因 AI 失效被弃用的周期缩短。

副量：峰兑份 b 。 一个人总价值中属于峰兑件的份额。必须诚实写明： **b 在现有任何记录体系里跑不动，因为字段不存在**——没有一个人事系统、一份简历标准、一张成绩单里有“经核验的事件兑付记录”这一栏。把这一栏补上，是本文第三层主张的落点（见第十节与第十五节）；在它被补上之前， b 只能在事件之后由结算处倒推，不能在期内由量具正读。这条“跑不动”不是本文的短处，它就是命题本身在记录体系上的投影。

七 五个场景

零情态词判据一句问到底：**在正式记录里，一个人扛住那次事故所动用的东西，被记在谁名下、算多少？**拿它去问五个真实场景，得到五个互不相同的答案——其中两个是查出来的，不是从命题推出来的。

1. **应届生的简历。** 一个学生大四那年顶住过一次实验室事故、把一个崩溃的小组项目从死线边上拖了回来。简历标准（无论哪家的字段规范）里没有任何一栏能

装这件事的核验记录；写进"自我评价"就与千篇一律的形容词混在一起。答案：**无字段**。

2. **民航驾驶舱（查出来的）**。航空业恰恰是给峰兑立了字段的少数行业：飞行小时逐条入册、机型资质、定期模拟机复训记录在案，且飞行员按资历小时计薪——不管这一年遇没遇到紧急情况。答案：**有字段，但字段记的是"在场时长"，不是"能扛多少"**；1500 小时之争的双方都承认，时长只是那份储备的粗代理。

3. **医院的备班（查出来的）**。值班备召（on-call）有明码补贴，常见按正常时薪的几分之一计。答案：**有字段，但它给"人被拴住的时间"定价，不给"被召到时他能接住什么"定价**——一位能接住主动脉夹层的备班医生与一位接不住的，备班费同额。

4. **电网**。容量拍卖给"随时可发"专门付费，兆瓦数事前可验（机组可以被点火测试）。答案：**有完整字段——因为它的容量可以不动用事件而被核验**；这恰好反衬出人身峰兑的难处：人没有"点火测试"，见第十三节。

5. **公司裁员名单**。一位平时产出平平、但过去三年两次在系统雪崩夜把公司捞回来的工程师，在按季度产出排序的优化名单上位置靠后。他被裁的那个季度，成本报表改善。答案：**曾有事件、无人建账；兑付记录随他离职一并蒸发**。

五个答案互不相同；第 2、3 条是查档得来的，而且它们不顺着命题走——它们证明字段可以存在，从而逼出更准的一刀：**现有字段给"在场"定价，没有一个字段给"承接能力"定价。缺的不是待命费，是承接力的核验与建账**。

八 量具是怎么死的：地板、差额、渐蠕

把机制串成一条线，用一个日常参照物：外卖平台上的商家评分。几年之内，几乎所有店都变成了 4.8 到 5.0——不是所有店都变好了，是评分这把尺子的整段量程被挤进了 0.2 分里。研究者在五个在线市场上实测过这条曲线并命名为**声誉膨胀**（Filippas、Horton 与 Golden）：评分者不愿伤害被评者，均值只升不降，系统"播下了自己失效的种子"。

毕业生评价正在走同一条曲线，但推力更硬。声誉膨胀靠的是评分者的心软；这里靠的是工具：AI 直接把每个人可提交产出的下限抬到接近上限。地板一抬，三步连锁——

第一步，差距缩进差额。量具没变，人群的分布被压扁；真实差异仍在（本文第十五节 F1 专测这一条），但缩进了量具分不出来的区间。此时量具的每一次出分都是一次无信息的裁决——非劣效性文献对此有现成的判词：没有灵敏度的比较，会在有效药和无效药之间读出等同。

第二步，等同结论渐蠕。每一届只与上一届比、每一次招聘只与上一批比，判等差额一次吞一点。教育侧的量具已经把这条曲线走完了：C 曾经是美国大学最常见的成绩，走了半个世纪，A 成了最常见的成绩。逐年看，每一年的评分都合格；合起来看，标准换了一个字母。

第三步，量具死亡外溢为量具弃用。当一把尺子对多数比较对失去分辨力，用尺子的人会扔掉它——这一步不是推论，是本文第十四节里已经验到的历史。

而且弃用不是终点，是循环的开始。把 2015 年以来毕业生筛选量具的更替排成一条年表，能看见周期在缩短：家庭作业式笔试与求职文书统治了十年，2023 年前后被生成工具一年内击穿；线上限时测评顶上，随即败给实时辅助；2024 到 2026 年，现场手写、口头答辩、结对实操轮番回潮，而每一种新形态从"有效"到"备考产业出攻略、AI 出模板"的间隔，已经从以年计缩到以月计。每一次更替都被宣传为"这次的量具读得出真本事"；本文的判断

是，只要它仍是平时量具——可预约、可排练、有重来按钮——它就仍在地板的射程之内，区别只是死得快慢。（其中“就本人产物做无脚本追问”这一支是同批《脱稿步》的领地，它能修复的份额与修复不了的份额，两篇的分界在第十二节。）

还有一条常见的自救要单独驳一下：改用强制排名。既然分数挤在一起分不开，那就规定必须排出一二三——排名制确实每次都交出一个序，但**有序不等于有信息**：当真实差距在差额之内，排出来的序在两批评委之间不可复现，序的内容是噪声加评委口味。这就是第六节手册把排名制列为“不适用”的原因：它不治量具死亡，只是禁止量具承认自己死了。

这里必须与两位近邻分开。**Spence (1973) 的混同均衡**说的是：当高低类型发出信号的成本趋同，信号失去分离力——他的解是造一种成本更高的新信号。本文说的比这多一层：在 AI 条件下，任何新造的平时信号都会在数月内被地板追平（家庭作业式笔试被生成工具追平、改现场笔试、再被实时辅助追平……），**量具的预期寿命本身在缩短**；分离不是找不到新信号，是“平时信号”这个类别整体在失效。**声誉膨胀**的病因在评分者的激励，处方是改激励（如私密评分）；本文的病因在被评物的地板，激励改得再好，尺子也读不出已缩进差额的差距——可裁决的分界：若把评分激励修好后分辨力恢复，本文错，声誉膨胀对；若激励修好而 m 照升，本文对。

九 欠付的推导：为什么按产出付钱必然漏掉峰兑

这一节不引任何外援，只用劳动市场自己的四步把结论推出来；推完之后再看到别处有没有人独立撞上过同一结局。

前提一（兑现分布）：一个人价值中存在兑现尖峰的成分——绝大部分兑现集中在稀发状态。这不是假设，是峰兑的定义（第四节第二条性质）。

前提二（计费的核验约束）：工资只能挂在双方都核验得了的量上；而由期内盲（第一条性质），事件到来之前唯一可共同核验的量就是平时产出。于是不管合同怎么写，可执行的那部分计费必然落在产出上——这不是雇主的选择，是核验约束下的唯一可行合约。

前提三（竞争定价）：可核验产出的价格由市场竞争压向其边际供给成本；AI 正把这块地板压得更低，并把更多岗位推向更纯的产出计费（零工化、按件外包、OKR）。

四步推论：一份价值若其兑现不走产出通道（前提一），而工资只能走产出通道（前提二、三），则持有它拿不到对价；维持它却有真实成本——训练时间、编制冗余、随时可召的机会成本。**无对价而有成本的东西，理性的持有者会减持；而被减持的恰是期内读数为零的成分，所以减持在每一期的账面上只表现为成本下降，不表现为任何损失。欠付、减产、账面改善，三件事由同一个推导一次给出，中间没有用到任何一个电学名词。**

推导站稳之后，再看旁证：电力经济学在自己的领域独立撞上并命名了同构的结局——只按发出的电量付费，则专为峰荷备着、常年不转的机组必然挣不回成本，业内称**缺钱问题**（missing money），奠基可回溯到 Boiteux (1949) 的峰荷定价。同构逐条可对：兑现尖峰对峰荷、产出计费对电量计费、竞争地板对电价上限。更要紧的是它给出了本推导的**制度解存在性证明**：容量市场——不为产出付费、专为“站在那里”付费的第二种合约——真的被发明出来并运转了二十年。电网不是本文的论据，是本文推导在另一颗行星上的既有落地；2021 年 2 月得克萨斯的大停电，则是那颗行星上欠付走到尽头样子。

有一句必须替产出计费说的公道话：这一步在它自己的范围里是**正确的**，不是失误。当价值的绝大部分确实平摊在常态里（AI 之前的多数白领工作正是如此），按产出计费与按总价值计费给出几乎相同的排序，而前者的计量成本低一个数量级——用便宜的账代理贵的账，是理性的制度设计。检验它是机制还是事故只需一问：把这一步“改对”（改为按总价值

计费)，空位会不会消失？不会——因为峰兑按定义无法被期内计量，总价值的账根本开不出来。改对了这一步空位仍在，所以它是机制；这也解释了为什么没有任何一方需要犯错，欠付就会发生。

电网的解法是发明第二种付费——容量拍卖。劳动市场有没有对应物？查下去会发现零星的雏形：律师的常年顾问费（retainer）就是一种人身容量费——但注意它成立的前提：**只付给已有事件记录的人**。没有打过硬仗案底的新律师，没人给付顾问费。这条观察把本文引到毕业生身上最锋利的那一点。

十 毕业生：零事件载体

前面九节讲的机制适用于所有劳动者；"毕业生"三个字加进来，问题陡然收紧一档。

一位干了十五年的工程师，就算平时量具对他也盲，他名下多少散落着几次可考的兑付：某年某夜的事故复盘里有他的名字，某个客户记得是谁把单子救回来的。这些记录粗糙、无字段、难核验，但存在。**应届毕业生是峰兑的零事件载体：他身上那份灭火器价值可能真实存在，但没有任何一次结算发生过，因此不可核验、不可定价、不可作保**。大学四年恰恰是被制度性地保护得不发生真实事件的四年——考试可以补考，项目可以延期，一切都有重来按钮；而有重来按钮的地方，按定义没有峰兑结算。

于是 AI 时代毕业生就业市场的三件怪事可以一句话贯通：**平时量具死了（所以简历与笔试比不出人），路径换算失锚了（所以学历与绩点被弃用），而事件记录为零（所以峰兑侧也无从定价）——三条读数通道同时断供，落在同一批人身上**。入门级岗位的收缩，通常被解释为"AI 替代了初级工作"；本文给出并列的第二重解释：初级岗位历来是市场为新人**垫资第一次事件**的机构（学徒制的现代残骸），当平时产出不再需要新人，垫资的商业理由消失，市场便停止为零事件载体发放"第一次结算"的入场券——这不只是岗位消失，是**建账通道消失**。而一个不再给新人记第一笔账的系统，十年后将在资深端发现自己无账可查——这与电网欠付备用容量、多年后在寒潮里发现无容量可调，是同一道算术。

所以"核心竞争力将是什么"的完整答案分两层。**存量层：峰兑本身**——那份只在关口显形的承接力，它是量具死亡之后唯一一测不死的成分，因为它的结算根本不走量具。**通道层：第一张可核验的兑付记录**——在一个建账通道萎缩的市场里，谁先让一次真实的、无重来按钮的结算落在自己名下（一次真实用户的线上事故、一次真金白银的交付关口、一次有外部对手方签收的救火），谁就先把自己从②③不可区分的盲区里捞出来。

"可核验的兑付记录"不是修辞，它有四条构成，缺一条就退化为简历上的又一句自夸：**其一，真实对手方**——事件的另一头站着一个利益不由你支配的人（客户、用户、事故的受损方），不是老师、不是队友；**其二，真实赌注**——事件当时有不可撤销的损失在场，搞砸了有人真的损失；**其三，无重来按钮**——当时不存在补考、延期、回滚的选项；**其四，对手方签收**——事后有那一方留下的、可被第三人查证的结算痕迹（一封客户的复盘邮件、一份事故报告里的名字、一笔实际发生的赔付或止损）。拿这四条去过一遍简历上的常见素材：课程大作业四条全无；有真实用户的兼职顶上一到三条；一次线上事故的处置四条可以全齐。给毕业生的操作性推论只有一句：**去有真实结算的地方，让至少一次结算写上你的名字；量具上的又一个满分，已经买不到它买过的东西了**。

这种账本不是空想，航空业早就记着一本。美国的航空安全报告系统（ASRS）自 1976 年起接收一线人员自愿提交的未遂事件报告，几十年攒下百万量级的"差一点"档案——**事件是可以建账的，连未遂的事件都可以**；只是那本账为安全归因而设，从不用于给报告人定价。存在性证明有了，缺的是把同类账本接到人的结算上，并跨过第十五节反噬预言里的那道坎。

最后把"入门岗为什么曾经存在"的账算清。培训经济学早就写明：学徒期是一份跨期合约——新人以低薪垫付自己的成长成本，企业以平时产出回收垫资。AI 抽走的正是合约里"新人的平时产出"这一项：当初级产出可以按边际成本从工具处买到，企业侧的回收项清零，合约自动流产。**所以入门岗的消失不该被读成"新人没用了"，该被读成"旧合约的对价没了"**——而新的对价（新人身上唯一买不到的东西）恰恰是尚无账可查的峰兑。市场缺的是让这笔对价可写进合约的记账法；这就是为什么本文把赌注的后半句押在字段上，而不押在情怀上。

十一 与最近的几种说法的分界

本站立在一圈近邻中间，每一位都必须点名，并给出可裁决的分离。

Weisbrod (1964) 的期权价值最早给"可用性本身"定了价：一家你从不去的镇医院对你有价值，因为它在那里。分离在于：Weisbrod 的定价走需求侧的支付意愿，前提是价值可被受益人事前估计；本文说的是人身承接力连**排序**都做不到——信息失灵先于金融失灵。判定形式：若 Weisbrod 成立而本文不成立，应能观察到对个人承接力（而非在场时间）运转的事前定价市场；实况是这类市场仅在既往事件记录处出现（第九节顾问费观察），与本文一致。

Cyert 与 March (1963) 的组织冗余把"松弛"当作组织吸收冲击的缓冲。分离：冗余是组织层的资源存量，本文是个人价值的构成份额加一条市场欠付机制；冗余论不解释为什么欠付在账面上读作改善。

Schmidt 与 Hunter (1998) 的甄选效度传统是应用领域里最硬的正主：一个世纪的元分析告诉招聘者哪种量具最能预测工作表现。本文的分离藏在它的效标里：这一传统所验证的"工作表现"，几乎全部是主管评定的**常态表现**——稀发事件因其稀发，进不了任何研究窗口的效标变量。所以一把在该传统下"高效度"的量具，效的是常态那一列；它对峰兑列的效度是未定义，不是已证明。判定形式：取效标为"事件表现"的验证研究（如空难机组表现），若平时量具在其上仍保持效度，本文错。

基于胜任力/技能的招聘运动（NACE 系的实务共同体）是弃用绩点之后实务界自己的答案：换更细的量具。本文的预测与它正面相撞：技能量具同样是平时量具，将以更快的周期被地板追平——2026 年的实况已见端倪，七成雇主自称转向技能招聘，同时抱怨求职材料全被 AI 写得千篇一律。谁对，看技能测评的更替周期。

行为面试与 STAR 法是本节必须放在正主席位上的一位——它在成稿最后一遍排查中才被认出来，而它其实是全场离本文最近的现役制度。八成五到九成的雇主在用"讲一次你处理危机的经历"（Tell me about a time...）这类问题选人，其方法学谱系可回溯到 Janz（1982）的行为描述面试，其信条一字不差地就是本文的半句话：过去的行为预测未来的表现——**劳动市场早就在试图读事件，这是对峰兑存在需求侧的最强证据，本文不能也不必否认它**。分离在结算结构上，共三刀。第一刀：STAR 读的是**本人自述、无对手方签收**的叙事——事件是否发生、损失是否真实、扛住的是不是他，全凭一张嘴；它是对事件的**转述**，而转述属于平时显露，可背诵、可教练、如今可由 AI 批量代写（求职辅导业已把 STAR 故事模板化成产业）。第二刀：正因为它是平时显露，它全额继承量具死亡——当每个候选人都带着三个结构完美的 STAR 故事进场，这把尺子的量程同样被挤进差额，本文预测它的分辨力衰减曲线与笔试同形、只是相位滞后。第三刀：对毕业生，STAR 预设了事件存量，于是备考产业教学生把课程作业包装成"危机"——零事件载体的窘境在民间话术里早已显形（市场把有事件的人叫"经过事的""battle-tested"，却拿不出任何一栏去核验这个形容词）。判定形式：给同批候选人配双盲事实核查（联系事件对手方核实），若核实前后录用排序不

变，说明 STAR 已在读真实峰兑、本文的"核验缺口"不成立；若排序显著翻动，本文对——**缺的从来不是问事件的意愿，是签收过的账。**

顺着这一刀补两位谱系上游。**Becker (1964)** 的培训经济学回答过"谁为新人成长期买单"：通用技能由学徒自付（低薪期），专用技能由企业分担——这是入门级岗位作为垫资机构的理论底座；本文第十节的"建账通道消失"是它在 AI 条件下的推论：当低薪学徒期不再能以平时产出回本，Becker 式合约的企业侧动机清零。**可携带凭证运动**（开放徽章、可验证学习与就业记录一类）是近年离"建账"最近的工程实践；分离一句：它核验并携带的是**胜任力凭证**——平时量具的产物——不是有对手方签收的事件结算；把它的管道接到事件结算上，恰是本文赌注后半句愿赌服输的那条路。

监管压力测试（2009 年银行业 SCAP 起）是方法学上离本文最近的发明：不等真实危机，主动造一个假想尾部去测。分离：压力测试之所以在银行业可行，是因为资产负债表可以在纸面上承受假想冲击而不动用真实损失；人身峰兑的构成条件（无重来按钮）恰恰不可纸面化——这条分离同时是本文第十三节第三处反向约束的正面。

Perrow《Normal Accidents》证明紧耦合系统里尾部事件不可根除；**Taleb《Antifragile》**主张冗余与凸性是对尾部的正确姿态。两位都在"尾部重要"这一侧，但都没有给出本文的两件：期内量具对承接者的盲，与产出计费对承接者的欠付算法。Taleb 的"切肤之痛"（skin in the game）离得最近，分离在方向：他要求决策者暴露于尾部损失（惩罚侧），本文要求结算体系看见尾部承接（建账侧）；一个防坏人，一个付好人。

声誉膨胀（Filippas-Horton-Golden）与 **Spence 混同**的分界已在第八节给出判定形式，此处不重复。**Arrow (1963)** 指出医疗质量不可观察时社会以执业许可制回应——许可制正是"用路径与考试代理不可观察质量"的始祖，本文第四节第二重否定的谱系上游。**Goodhart 与 Campbell** 的测量反噬与本文正交：他们讲量具被压力扭曲，本文讲量具无压力也死于地板；判据是撤掉考核压力后分辨力是否恢复。

以上诸位从七八个方向逼近同一处，而各差最后一步：有人看见了可用性有价（Weisbrod、Taleb），有人看见了量具会死（声誉膨胀、Spence），有人看见了尾部进不了效标（甄选效度传统自己的窗口限制），有人发明了给待命付费（容量市场、备班费）。**没有一位把四件事钉在同一个存在物上：期内盲、事件结算、不可点播、欠付即减产。**本文补的是这一钉。

划完界再退一步取一次交集：这一整圈近邻——心理测量、执照制度、压力测试、模拟机、测评产业——共同站着一块更深的地：**一份能力的证据，可以在不动用它所防备的事态的情况下取得。**整个评价文明建立在"可以不烧真火而考出救火"的信念上。本文只对一份额否认它：对峰兑那一份，证据与事态不可分离；这就是为什么它测不死，也是为什么它建不了平时的账。

十二 与同批六篇的分界

同一道题在同一天里已被回答过六次，六篇与本文同源同批，逐一交代，并各给一条决定性分界。

《敞问》答：竞争力是只能由本人合上的自我之问。它的结算人是本人，本文的结算人是事件——外部世界那次没有商量的关口。分界：一个人可以敞问满怀而从未被任何事件结算（③④不可分），也可以敞问寥寥而兑付累累。

《现配优势》答：优势是每次交手当场做出的配组，每一回合都在读它。本文恰好补它的余集：有一份价值连"当场"都不读——平常回合无论多少次，对峰兑全盲。分界观察：在

配组完全不变的常规回合序列里，若某人离开后组织的常规表现不变而尾部事件成本骤升，则存在现配读数从未捕捉的份额，本文对。

《零判份》答：合格产出里有判别力为零的一份，因为世界上等价物太多。零在产出与世界之间；本文的零在量具与人群之间——差异真实存在，仪器读不出。判决性检验：对量具上并列的候选人做专家事后裁决，若裁出真实差距，本文对；若裁不出（真等价），零判份对。

《脱稿步》答：竞争力显形于档案延长线上被发起的一次追问。追问是可以在任何一个平常日子被主动发起的廉价探针——它能修复的，正是本文宣布不可修复的那部分吗？分界写成让步：追问能考出程序性与理解性的份额，那部分不属于峰兑；峰兑的定义收窄为追问也够不到的份额——真实损失承受，其构成条件（无重来按钮）使任何被发起的探测自动降格为模拟。若有一种平时探针能复现事件排序（跨场次秩相关 ≥ 0.7 且两年不失效），本文第一层判伪。

《常席差》答：竞争力是人作为跨配组唯一不变项时才读得出的残差——仍然是常规回合的读数，只是读法更聪明。分界：常席差需要多回合轮换记录才显影；峰兑在任意多的常规回合里都不显影，只认事件。两者若在同一台账上给出相反评价，看该台账里有没有一次真实结算：有，以结算为准。

《回改》答：已作出的裁定会被后到证据翻改。回改的前提是期内曾有裁定可翻；峰兑的份额上期从内从无裁定——无可翻之账。两篇是互补件：回改管“账错了”，本文管“没建账”。

同产线稍早的《兜底份》（论承载者的保守预算覆盖数）是站内血缘最近的一篇：它测承载者在最坏情形预算中覆盖了几成——供给侧的备灾充足率，可从预算表事前算出。本文测的是需求侧与计量侧的失灵：覆盖得再足，期内量具照样读零、市场照样欠付。分界观察：一个兜底覆盖数很高的组织，其承载者在绩效台账上仍不可辨——两个读数一个事前可算、一个只能事后结算，不同量。

按归档口径先记一笔：本篇落在「什么之问 × 显露待解」一格的逆解侧，与同批《现配优势》《零判份》《脱稿步》《常席差》《回改》同格，判决性分界已逐条给出。六篇合起来有一句必须坦白的話：同一道题七次降落，七个 Z 全部落在“读数与价值错位”这片大陆上——这本身是一条关于这道题的读数：**AI 时代把“竞争力”从一个能力问题改写成了一个计量问题**；七篇是同一座冰山露出的七个角。本篇的角在计量死角的最深处：其余六篇至少还相信有某种读法（配组、残差、追问、回改、自问）能把价值读出来，本篇指认了一份任何读法都读不出、只能等结算的余额。

十三 三处反向约束

三处真实的削弱，各改掉本文的一句话。

第一处：保险业。寿险、职业责任险早就在给人身峰时事件定价——“峰时不可定价”的强版说法必须删除。收窄后的准确句：保险给**损失**定价，不给**承载力**定价；它算的是你出事赔多少，不是你能替系统扛住多少。前者有精算表，后者正是本文说的无账户。

第二处：强管制安全行业。民航、核电、消防、军队整建制地给人身容量立了字段（时数、演训、编制、战备等级）——“制度必然欠付”必须收窄为“**产出计费型市场必然欠付**”。这些行业的共同点是容量条款由管制强加，不由市场自发生成；而 AI 时代的用工形态正把更多岗位推离管制、推向纯产出计费——适用域的边界本身在朝不利于劳动者的方向移动。

第三处：全动模拟机。航空模拟训练对紧急程序的效度有几十年证据，"模拟考不出峰兑"的全称版必须删。收窄后的句子：模拟能核验**程序层**（操作序列、决断树），核不动**承受层**（真实损失下的执行完整度）——1500 小时之争恰恰卡在两派对这条缝隙宽度的估计不同。本文只对承受层主张期内盲；程序层请交还给模拟机与《脱稿步》。

十四 事先登记的检验

本文第三层有一条不依赖峰兑概念的经验主张：**在一把有界量具上，顶格份额越过多数化阈值（量具饱和）之后，会出现市场对该量具的持续弃用；饱和在先，弃用在后。**检验按事先注册执行：判负条款、不利结果处置、停止规则三件在读取任何数据之前写死存档，SHA-256 为 5cbea97092661389a89e42e3234a82efa00d9c5f39b5887040e309c7dac27937；注册中同时申报了作者对量级的既有印象，判负条款集中于事先不知道的三处（最新值、下滑起点、先后次序）。

量具取美国大学成绩（A-F 有界量表），弃用方取雇主。数据两族：Rojstaczer 与 Healy 的全国成绩份额年表（400 余校、四百万在校生）；NACE 历年 Job Outlook 的"按绩点筛选候选人的雇主占比"。

P1（饱和）：支持。A 份额 1960 年约 15%，1988 年约 31%，2008 年达 43%，约 1998 年起 A 成为最常见成绩；43% ≥ 40% 的判负线。**P2（弃用）：支持。**雇主绩点筛选占比自 2019 年峰值 73.3% 降至 2023 年最低点 37%，2026 年为 42.1%；峰值至最新降幅 31.2 个百分点，远超 15 个百分点的判负线。**P3（次序）：支持。**A 份额越过 40% 的年份按两个见档点（1988=31%，2008=43%）线性内插约为 2003 年，保守上界为 2008 年（该年 43% 已见档）；均早于 2019 年开始的连续下滑。

三处不利侧写，照登。其一，弃用不是单调的：2023 年 37% 之后回升至 2025 年 46%、2026 年 42.1%——弃用曲线与劳动市场松紧同向波动（新人过剩的年份筛子会被捡回来），这提示"弃用"至少部分由市场松紧驱动，不全由饱和驱动。其二，越线年与弃用起点之间隔了至少 11 年，如此长的滞后使"先后一致"不能升格为因果——按注册约定，第三层的措辞钉死在"次序与机制预测一致"，不得多走一步；且弃用的陡降段与疫情期重叠，存在明显的混杂事件。其三，全国成绩系列最近的扎实档位停在 2008-2013 年段（其后仅有校级零星数据，如个别名校 A 份额已至七成九）——**给全国量具饱和度和建账的那个国家级仪表盘，自己也停更了；这件事本身与本文命题同形，如实登记。**

按注册约定的处置：三条均获支持，第三层主张保留原措辞；不利侧写一并入文，机制归因维持在"先后一致"档。

十五 证伪条款、赌注与反噬

分层引用声明：本文可被分层引用——第一层（峰兑作为存在物）倒了，第二层（辨别格与盲读率）仍可用；第二层倒了，第三层（饱和先于弃用的经验次序，第十四节）独立成立。

证伪条款，只列可执行档：

- **F1 (差异真实性)**：取任一场 $m > 0.6$ 的筛选场次，对量具并列的候选人做双盲专家事后裁决 ($n \geq 30$ 对，裁决者不见量具分)。若专家也裁不出稳定差距（复裁一致率与随机无异），则“差距真实存在、仪器读不出”被驳，本文塌向《零月份》。
- **F2 (盲读率走向)**：按第六节手册连续三年测同类场次的 m 。若 AI 渗透率上升期 m 不升反降，第八节机制判伪。差额取法、粒度与第六节同一口径。
- **F3 (量具寿命)**：编码 2015–2030 年主流毕业生笔试/作业形式的采用一弃用年表。若量具中位存活期不缩短，第八节“量具死亡加速”判伪。
- **F4 (欠付一减产)**：在可获得事件台账的行业（民航、电网运维、SRE），检验“产出计费纯度越高的组织，人身备灾配置越薄、且薄化当期常规指标越好看”。三个变量各有现成代理（计薪结构、编制冗余率、季度成本），方向若反，第九节的推导被驳。
- **F5 (建账通道)**：追踪入门级岗位收缩与“新人首次真实结算机会”的关系。若新人首事件年龄不后移，第十节第二重解释判伪。
- **F6 (杀死条款)**：若出现一种平时量具，在 AI 饱和条件下对毕业生的排序跨场次秩相关 ≥ 0.7 且两年不被地板追平，本文第一层判伪，全文降格为该量具出现之前的历史注脚。

赌注 (写死日期)：到 2030 年 12 月 31 日，NACE 口径的雇主绩点筛选占比不会回到 60% 以上（量具死亡在差额之内不可逆）；且主流简历/人事数据标准中不会出现被前一百大雇主中至少十家采用的“经第三方核验的事件兑付记录”结构化字段。前半句错，第八节错；后半句若错——本文乐于错，因为那意味着建账通道被修好了，而这正是本文的处方。

反噬预言，明写：峰兑一旦进考核，最短的达标路径是制造事件——救火英雄文化的病理档案里早有先例，包括为了扑火而放火的人。部分缓解是把结算权交给外部对手方（事件须有真实受损方签收，不得由本人或本组织自证）；但只要“有事件”优于“无事件”，让小火焰多烧一会儿再扑的激励就删不干净。**我看不到完整的出口，只看到结算权外置这半个。**

伦理三条（本读数属可通约量纲，自带此件）：① 盲读率与峰兑份指的是位置与制度，不是人；不得据以断言任何个体“没有峰兑”。② 不得为了改善任何一方的峰兑读数而制造、放任或延迟处置风险；安全关键系统中不得为制造“事件机会”安排新人顶岗。③ 凡已按本读数作出不利人事安排的，须公开差额取法与编码规则，并提供复核通道。三条缺一，本文不可被引用于任何考核场景。

十六 自反与结语

用本文自己的判据给本文定位：这是一份典型的峰兑式产出。它的全部主张关于稀发事件，因此没有一条能在发表当期被常规读数验证；它平时的“绩效”——被引、被赞、被转发——恰恰是它自己宣布无信息的那类量具；它真正的结算日在 2030 年的赌注与某一次尚未发生的、大面积的人才断供之间。**引用本文的人持有它的峰兑风险：如果那一天不来，本文一文不值，而且在此之前的每一天它都显得言之凿凿。**这个定位不是谦辞，是本文对自己命题的第一次执行。

结语收在毕业生身上。这一代人拿到的是三条同时断掉的读数通道：量具读不出你，学历换不出你，事件账上还没有你。前两条断供不是你的错，也轮不到你修；第三条是唯一在你手里的。灭火器的价值从来不在压力表上，在火里。去有真实结算的地方，让第一笔账落在你名下——那是量具死透之后，这个世界还剩下的、唯一读得出你的方式。

附录一 盲读率清点空表与边界例

场次编号 | 职位 | n | $C(n,2)$ | 主量具 | 判等差额 (取法注明: 基线复现最小分差/2×复评标准差) | 盲读对数 | m | 剔除数及原因

边界例存档: 总分差恰等于差额=分得开; 同分并列=盲读对; 跨轮候选不拼对; 一场只认一把主量具。差额一经取定, 正文第六节、第十五节 F2、赌注三处同值, 改任一处即全篇作废重算。

附录二 预注册回放

预注册全文另附随文文件 (prereg_run7.md), SHA-256 同第十四节。三条判负条款 (P1 饱和线 40%/35%; P2 弃用降幅 15pp/5pp; P3 越线年早于下滑起点年)、不利结果处置与停止规则均写于取数之前; 既有量级印象已申报。复算脚本四行算术: 两个见档点 (1988=31%, 2008=43%) 线性内插得越线年约 2003; NACE 峰值 73.3% 减最新 42.1% 得降幅 31.2pp。任何人可用同两族公开来源复算。

附录三 兑付记录字段规范 v0.1 (第十节四条构成的可执行版)

字段	内容	由谁写入	核验方式
事件标识	时间、地点/系统、事件一句描述	承接人发起	与对手方记录比对
对手方	利益不由承接人支配的另一方 (客户/用户/受损方) 及其联系凭证	对手方确认	第三人可致询
赌注	事发当时在场的不可撤销损失 (金额/停摆时长/伤害面)	对手方申报	账目或工单佐证
无重来声明	当时不存在补考/回滚/延期选项的说明	双方共签	与系统日志对勘
承接动作	承接人实际做了什么 (动作清单, 非形容词)	承接人写、对手方核	逐条可指认
结算结果	实际止损/兑付了多少	对手方签收	与赌注字段对减
签收凭证	对手方留下的可查证痕迹 (复盘邮件/事故报告/赔付单号)	对手方	第三人可复核

签收协议三条: ①任何字段不得由承接人单方自证; ②本组织内部奖状不算对手方签收 (利益同向); ③记录一经签收即不可由任一单方修改, 改动须双方重签并留痕。此表即本文赌注后半句所押的那个"结构化字段"的最小可行形状; 它同时是第十五节反噬预言的第一道闸——签收权外置, 正是从字段设计处执行的。

注释与材料来源 替加环素: Prasad P., Sun J., Danner R.L. 等, "Excess Deaths Associated With Tigecycline After Approval Based on Noninferiority Trials," *Clinical Infectious Diseases* 54 (2012): 1699-1709; FDA Drug Safety Communication, 2013-09-27 (黑框警告)。灵敏度与渐蠕: ICH E10; FDA, *Non-Inferiority Clinical Trials to Establish Effectiveness* (2016)。时数之争: Public Law 111-216 (2010, "1500 小时规则"); ICAO MPL 框架; NTSB, Colgan Air 3407 事故报告; Sullenberger 各次国会证词。容量市场: Boiteux M. (1949) 峰荷定价; Cramton P., Ockenfels A., Stoft S., "Capacity Market Fundamentals," *Economics of Energy & Environmental Policy* (2013); 2021 年得州停电各调查报告。声誉膨胀: Filippas A., Horton J.J., Golden J.M., "Reputation Inflation," *Marketing Science* 41(4) (2022): 733-745。甄选效度: Schmidt F.L., Hunter J.E., *Psychological Bulletin* 124 (1998): 262-274。期权价值: Weisbrod B.A., *Quarterly Journal of Economics*

78 (1964): 471-477。行为面试: Janz T., *Journal of Applied Psychology* 67 (1982): 577-580; 各校就业指导中心 STAR 法通行教程与雇主使用率调查 (85-91%)。培训合约: Becker G.S., *Human Capital* (1964)。未遂事件账本: NASA Aviation Safety Reporting System (1976-)。混同均衡: Spence M., *QJE* 87 (1973): 355-374。组织冗余: Cyert R.M., March J.G., *A Behavioral Theory of the Firm* (1963)。质量不可观察与执照: Arrow K.J., *American Economic Review* 53 (1963): 941-973。事故与尾部: Perrow C., *Normal Accidents* (1984); Taleb N.N., *Antifragile* (2012)。成绩份额: Rojstaczer S., Healy C., "Where A Is Ordinary," *Teachers College Record* 114 (2012); gradeinflation.com (2016 年更新)。绩点筛选: NACE Job Outlook 2019-2026 各期 (73.3%→37%→46%→42.1%)。一分一段: 河南省教育考试院《2025 年普通高招考考生文化成绩分数段统计表 (物理类)》官方锚点 (新华网河南频道 2025-06-26 转发)。预注册与复算脚本见文末哈希与随文文件。

字数与版本 全文实数 14,757 汉字 (机检口径, 含摘要、附录与文末件); v3.0 (V1/V2 同日作废并入; v3 完成欠付推导自立化、盲读密度实算、字段规范), 2026-09-02。本文一趟写成, 未经外部评审; 文中一切分数性自评均不存在——按写作纪律, 作者不为自己参与的文本发认证分。

人机分工声明 选题与题型由委托人给定; 机制构造、材料检索、预注册设计、数据复算与全文写作由 AI 完成; 替加环素、NACE、成绩份额等全部经验材料均取自公开可复核来源, 检索与复算过程存档。

第十二章 零判份

【章首】本章给可读条件立第三条：**相对世界样本**。一次合格产出可分成两份：世界还拿不出可复现同类物的那份（未见份），与判别力为零的那份（零判份）——后者不是低质量、不是抄袭，可以是全篇最漂亮、最费力的一段，它判别力为零只因世界已能拿出同类样本。读数是未见率 ρ ，且本章的最强主张是： ρ 的收窄由暴露时长驱动而非使用强度驱动——保密、限流、不发布都拦不住。它与《现配优势》的判决对照写在两章正文里：取同人同配组五年，表现稳定而判别力单调下降则本章对，两者同步则被彼章吸收。本章的预注册检验结果不利（一项仅获定性支持、两项未获裁决），已照原样写入并因此改写了一处核心表述——此处诚实是本章可信度的一部分。本章欠的账：保密条件下的同质多案例检验（ $N \geq 6$ ），未做。

大白话：外面到处买得到的那部分，不算你的本事

一道菜端上来，很漂亮，做了三个小时。但如果这道菜楼下三家店都做得出来，那么这三个小时里，有多少是“你的本事”？

《零判份》的回答很不留情：一份产出可以切成两半——世界还拿不出可复现同类物的那一份，和世界已经能拿出来的那一份。后者的判别力是零。

零判份不是“做得差的部分”，也不是“抄的部分”。它可以是全文最漂亮、最费力、你最得意的那一段——它判别力为零的唯一原因是：世界已经能拿出同类的东西了。

为什么今天要紧？

因为“世界能拿出来的东西”这个集合，正在以人类历史上从未有过的速度膨胀。昨天还很稀缺的一份分析、一段代码、一张设计图、一篇文案，今天变成了任何人三分钟可得的东西。

于是每个人的产出里，零判份的比例都在被动上升——你什么都没做错，你的判别力在自己下降。

这一章最锋利的一句话：这个差有半衰期。

它的收窄由**暴露时长**驱动，而不是由使用强度驱动。也就是说，一个东西只要在这个世界上存在得够久，哪怕你从不发布、严格保密、限制流传，它的判别力也会随着世界从别的路径长出同类物而收窄。

保密不能阻止它。这一点非常反直觉，但它把很多人的策略判死了：**囤积不是资产保值的办法**，因为收窄不需要你泄露。

读数是什么？

叫**未见率 ρ** ：你这份产出里，世界目前还拿不出可复现同类物的份额。

它的用法比听起来实际。你在动手之前问一句：**这件事我做出来之后，外面有没有现成的？**有，那么它属于零判份，做它的理由只能是“我需要这个东西”，不能是“这能证明我”。它不能证明你，一天也不能。

一个小场景。

一个设计师花了两周做了一套非常精致的品牌视觉。做完那天他很有成就感。三个月后他发现，同样风格的东西，别人用工具一小时能出十套。

他生气的点通常是“这不公平”。但按这一章的读法，事情不是公平不公平：他那两周的产出里，**未见率原本就不高**——他做的是一个世界已经能大量供给的东西，只是当时供给还没这么快而已。ρ 收窄不是有人抢了他，是暴露时间到了。

他真正该做的调整，是在动手前先问那句话，然后把力气挪到那些**世界还拿不出来的**部分上——比如只有他知道的这个客户的真实约束，比如需要他在现场跟人吵一架才能定下来的取舍。

它不是什么？

它不是“要做别人没做过的事”这种口号。口号是空的，这一章给的是一个**可以在动手之前查的动作**：先去看世界现存的样本，再决定自己的力气往哪放。

它也不是说漂亮和费力没价值——它们有价值，只是那份价值不用来判别你。

再看两个身边的例子。

手写的春联。三十年前，一手好字在村里是稀缺的，写春联的人受人尊敬。后来印刷的春联五块钱一副，做得比手写的还整齐。那手好字变差了吗？没有。**是世界能拿出来的样本变了**。这件事让人不舒服，但它和公平无关：判别力从来就是相对于世界现存样本而言的。

翻译。十五年前能把一份英文说明书译得通顺，是一项实打实的本事。今天这件事的未见率接近于零。而在同一个领域里，另一些部分未见率依然很高：把一份合同译得让双方律师都认可、在谈判现场把一句带刺的话译得既准确又不点火——这些世界还拿不出来。**同一个职业内部，零判份和未见份是混在一起的**，看不清这条线的人，会觉得整个行业塌了。

三种常见的误读。

第一种：以为它在说“越新奇越好”。不是。未见份不等于新奇，它是“世界拿不出可复现的同类物”。你手上关于这个具体客户的、这个具体现场的、这次具体谈判的东西，一点也不新奇，但世界拿不出来。**很多人在追新奇，而把身边真正未见的东西白白扔掉了**。

第二种：以为保密能保住它。书里明确判掉了这条路：收窄由暴露时长驱动，不由使用强度驱动。世界会从别的路径长出同类物，你藏与不藏都一样；藏着的代价是你还失去了它可能带来的全部反馈。

第三种：以为费力等于有价值。这是最伤人的一条，因为它牵扯到情感。你花了两周、做得很漂亮、自己也满意——这些都是真的，但它们与判别力无关。**这一章不否认你的辛苦，它只是说这份辛苦不该被拿来当作“我和别人不一样”的证据**。辛苦的价值在别处：它可能喂了你的手（上一章），可能换来了一个真实的交付。

给自己做个小检查。拿出你最近最得意的一份产出，逐段问：“这一段，外面现在拿不拿得出来？”诚实回答之后，看看剩下多少。剩下的那部分，就是你下一次该多花时间的地方。

它和旁边几个概念的区别。

和《**现配优势**》比的对象不同：一个是人对人（照单执行能复现多少），一个是你为世界（世界拿不拿得出同类物）。

和《**峰兑**》分属供给侧与测量侧，上面刚讲过。

和《**入账之前的技能**》有一处很实用的咬合：那一章告诉你，凡写进文件载体的东西都会进入可复制区；这一章告诉你，进入可复制区之后，它的判别力会随暴露时长单调收窄。

两章合起来是一条完整的因果链：入账→可复制→世界样本增加→ ρ 收窄。看懂这条链的人，就不会再指望靠“写得漂亮”来长期保值。

和《底流》的关系也值得记一笔： ρ 告诉你力气该往哪儿花（花在世界拿不出的那部分）， β 告诉你力气够不够（还有没有在被喂）。一个人可以 ρ 很高但 β 很低——想得很准，手却生了。

一句话记住它：这一章管的是**跟谁比**——不是跟同事比，是跟整个世界现在能拿出来的东西比。

怎么长出来：方法、技术与流程

个人层面（一个可以每周做的动作）。

先查世界样本，再动手。任何一个你打算投入超过一天的产出，动手前花二十分钟去看：这件事外面已经有什么？有多少？多好？拿得到吗？

查完之后做一个非常具体的划分：这次产出里，哪一部分外面已经有了（那部分尽管用最快的方法完成，包括用工具），哪一部分外面还没有（那部分自己亲手做，慢一点也值）。

这个动作同时解决了两件事：它让你的 ρ 变高，也顺便让你的 β （上一章的在喂率）花在了对的地方。

大学发生学教育：把“查世界样本”变成课程的固定工序。

现行教育有一个隐蔽的问题：绝大多数作业题，答案在世界上早已存在，且学生也知道。于是学生做作业时非常清楚自己在做一件零判份的事——**这正是“作业没意义”这种感受的真实来源**，它不是懒，是判断准确。

- **选题前的世界样本检索（必修工序）。**任何项目、论文、设计的选题环节，强制增加一步：提交一份“世界现存样本清单”——外面已经有什么，本项目与它们的差在哪里。这一步写不出差的题目，允许做，但要标明它是练习而非产出。
- **零判份声明。**提交作业时，学生自行标注：本作业中哪些部分是世界已能大量供给的（可以用工具、可以引用、可以直接拿来），哪些部分是自己新增的。**教师主要评后者。**这一条把“用不用 AI”这个已经无法监管的问题，转成了一个可以申报、可以讨论的问题。
- **未见率意识训练。**在低年级就要讲清楚：漂亮、费力、耗时，都不构成判别力；判别力只来自世界还拿不出的那一份。这句话越早说，学生越不容易把四年花在打磨零判份上。
- **课程题目的定期更新义务。**教研室应定期检查：本课程的作业题，是否已经可以被工具在几分钟内完成？如果是，这道题作为**练习**仍可保留（练习的价值在喂手，见上一章），但不得再作为**判别**依据。这两种用途必须分开，现在绝大多数课程把它们混在一起。

教师的动作。

把“这题外面能不能直接找到答案”当作出题时的一个显式变量，并且**如实告诉学生**这道题属于哪一类：是练手的（世界已有，但你得自己走一遍），还是判别的（世界还没有）。学生对被当作傻子非常敏感，而对被如实相告非常配合。

最容易做坏的地方。

把 ρ 理解成“要保密”。这一章明确判掉了这条路：**收窄不因保密而停**。围着不发的东西照样会贬值，而且贬值期间你还失去了它可能带来的所有反馈。

原文

零判份

论一次合格产出中判别力为零的那一部分，及其对"毕业生竞争力"的重估

副标题所承载的可裁决主张：一个人的一次产出，可以被分成判别力为正的一份与判别力为零的一份；现行的全部能力账本（学分、绩点、作品集、简历、面试评价、KPI）只记总量，不设这两份的字段；而在生成式模型普及之后，第二份在总量中的占比正在单调上升，且它的上升**不是由使用强度驱动的，是由暴露时长驱动的**——因此保密、限流、加密、"不公开发布"都不能减缓它。

摘要

问题：当一份合格的产出可以由机器在几秒内做出同类物时，"这个毕业生有能力"这句话还剩下什么可判定的内容？现行的三种回答互相排斥：看产出（成品合格即能力存在）、看过程（形成路径可核才算他的）、看位置（相对稀缺才有价值）。三种回答各自自洽，合起来无法并存。

承重命题：AI 时代大学毕业生的核心竞争力，不是他掌握的技能存量，也不是他能力形成路径的可核性，也不是他相对于同侪的稀缺位置，而是**他的一次产出中"这个世界还拿不出一个可复现的同类样本"的那一部分**——本文称之为**未见份**。与之互补的那一部分，在所有账本上被记为他的贡献而在判别意义上为零，本文称之为**零判份**。零判份不是低质量的那一份，也不是抄来的那一份；它可以是全篇最漂亮、最原创、最费力的一段。它之所以判别力为零，只因为世界已经能拿出同类样本。

材料与方法：三个互不相干的领域被拿来互相顶撞——结构化学的物质同一性判据、反兴奋剂科学的成绩归属判据、抗菌药物管理的药效有效性判据。推翻三者共同前提的那件材料取自结构化学自己的实证记录（利托那韦晶型事件）。一条证伪条款用公开数据做了真实检验，结果不利，已照原样写入第十一节并改写了本文的一处核心表述。

零情态词判据：*把这个人的这一次产出交给一个不认识他的人，那个人要花多久才能从公开可得的东西里拿出一个功能等价的样本？*

结论：竞争力不是一个人身上的属性，也不是他与他人之间的位差，而是**他的产出与世界现存可复现样本之间的差**。这个差随时间单调收窄，且收窄不因保密而停止。因此"培养核心竞争力"这句话在现行教育账本里是无法执行的——账本没有它的字段。

关键词：零判份 · 未见份 · 未见率 · 判别力消耗 · 能力账本 · 晶种场

English Abstract

Title: The Null-Discriminating Share: On the Portion of a Qualified Output That Carries Zero Discriminating Power, and a Reassessment of "Graduate Competitiveness"

Abstract: When a competent output can be matched within seconds by a machine, what remains judgeable in the claim "this graduate is capable"? Three prevailing answers are mutually exclusive: read the product (a passing artefact proves the capability), read the path (only a verifiable formation history assigns the capability to a person), read the position (only relative scarcity confers value). This paper argues that all three share an unstated premise — that

capability is the kind of thing that can be attributed to a person — and overturns it using structural chemistry's own empirical record (the ritonavir disappearing-polymorph episode of 1998, in which an identical molecule passing every identity assay ceased to be the same drug, and in which the original form became unobtainable in any facility where the new form had once nucleated).

The proposal: a graduate's core competitiveness in the AI era is neither a stock of skills, nor the auditability of its formation, nor a positional advantage, but the portion of a given output for which the world cannot yet produce a reproducible equivalent — here called the **unseen share**. Its complement, recorded in every ledger as the person's contribution while carrying zero discriminating power, is the **null-discriminating share**. The null-discriminating share is not the low-quality part, nor the plagiarised part; it may be the most original and most laborious passage in the work.

A pre-registered empirical check was run against public data. **The pre-registered source returned only qualitative results; the run was less informative than expected, and this adverse outcome is reported verbatim in §11, where it forced a revision of the paper's decay claim.** Secondary (non-pre-registered) evidence from a 2026 study of 60 language-model benchmarks indicates that discriminating power decays with **elapsed exposure**, not with **intensity of use**, and that private held-out test sets do not slow the decay — a finding that separates this account from both credential inflation and Goodhart-type gaming.

Keywords: null-discriminating share; unseen share; unseen rate; decay of discriminating power; capability ledger; seeding field

一 • 引言：一个不再有人敢问出口的问题

一个 2026 年毕业的学生，交出一份三万字的行业分析报告。数据翔实，结构清楚，结论有依据，没有一处引用错误。

十年前，这份报告能证明一件事：这个人会做这件事。

现在，它证明不了。不是因为报告变差了，是因为这句“能证明”依赖的东西悄悄消失了：过去，一份这样的报告在世界上是稀有的，稀有到“它出现在你手上”这件事本身构成信息。现在不是了。

于是三条路被同时提出来，各自都有人在认真走：

第一条：不管过程，看成品。 招聘方给候选人一道真实任务，四小时，随使用什么工具。做得好就是好。这条路的哲学是干净的：**一样东西显露出来的样子，就是它的全部身份。**

第二条：成品不够，要看路径。 于是有了闭卷、监考、过程留痕、草稿版本历史、口头答辩、AI 使用检测。这条路的哲学同样干净：**同一个成品，由不同的路径产生，就是两样不同的东西。**

第三条：成品与路径都不够，要看它在什么场里成立。 这条路说：一个技能有没有价值，取决于此刻还有多少人（多少台机器）具备它。这是最冷的一条，也是最被回避的一条。

这三条路互相取消。第一条说路径是脚手架，拆掉不留痕；第二条说路径正是身份本身；第三条说这两条争的东西都由外部条件场裁定，跟这个人无关。

本文不打算在三条里选一条，也不打算调和它们。本文要做的是把它们共同踩着的那块地翻过来。

路线图：第二至四节分别把三条路推到它们各自最硬的形态；第五节说出它们共同假定了什么；第六节给出推翻这个假定的那件材料——它来自三条路之一自己的实证记录；第七至九节给出由此长出来的那个东西、它的辨别装置与读数；第十节给出清点手册；第十一节报告一次真实的数据检验及其不利结果；第十二节与最近的既有说法逐条分界；第十三节写出证伪条件；第十四节是本文对自己的定位；第十五节是一条写死日期的赌注；第十六节写出本文不适用的边界与伦理约束。

二 · 第一条路：成品即身份（结构化学的读法）

化学在一百九十多年前解决过一次同样形状的问题。

1828 年，维勒从氰酸铵做出了尿素。这件事的意义不在于“人造了有机物”，而在于它确立了一条此后从未被正面挑战的判据：**做出来的那个东西，与狗尿里析出的那个东西，是同一东西**——因为它们的性质完全一致。至于它是从活体里来的还是从坩埚里来的，不进入身份判定。

这条判据在现代结构化学里被推到了极致。一个化合物的身份由一组可测量的东西给定：分子式、连接方式、立体构型；确证靠一组独立的谱学证据互相印证——核磁、质谱、红外、单晶衍射。合成路线写在论文的实验部分，但它不进入身份。**同一个分子，二十条不同的合成路线做出来的，是同一个分子。**谁做的、怎么做的、做了几次才成功、中间报废了多少批，全部是脚手架。脚手架拆掉，楼还是那座楼。

把这条判据搬到毕业生身上，得到的正是第一条路：**看成品**。报告好就是好。用了 AI？就像用了自动柱层析仪、用了商业化的起始原料、用了别人发表过的反应条件——化学从不因此判定产物不是产物。

这条路的力量在于它对“过程审查”有一个极难反驳的诘问：**你凭什么认为路径进入身份？**如果两份报告在任何可测量的维度上都无法区分，坚持说“这一份是他的能力、那一份不是”，你依据的是什么可测量的东西？

这一家自己知道、但它不往这个方向用的一个事实，留到第六节。

它把什么当成单独够用的那一样：显露。

三 · 第二条路：路径裁定身份（反兴奋剂科学的读法）

竞技体育面对的是同一个问题的极端形态，而且它已经面对了六十年。

一个成绩是一个数。9 秒 58 就是 9 秒 58，秒表不知道这个人昨天吃了什么。如果按化学的判据，服药跑出的 9 秒 58 与干净跑出的 9 秒 58 是同一东西。

体育不接受这个结论。它花了六十年建立一整套判据体系，其唯一目的就是**让形成路径进入成绩的身份**。

而这套体系最值得注意的地方，不是它检测得多准，而是它在什么地方**放弃了**。

早期的反兴奋剂是"查物质"：列一张禁用清单，测样本里有没有。这条路失败得很彻底，原因是结构性的——检测方法必须先公开或至少先存在，而规避方法总是后到。清单是静态的，规避是动态的。到了 2000 年代，这个不对称已经无法靠加大投入弥补。

于是有了**运动员生物护照**。它的做法是一次判据的迁移：不再问"你体内有没有 X"，改为问"**你自己的纵向基线有没有出现无法由生理解释的断裂**"。判据的锚点从一张外部清单，挪到了这个人自己的历史。

这个转向承认了一件重要的事：**当外部的绝对判据可以被适应，唯一还站得住的判据是个体自身的连续性。**

搬到毕业生身上，这是第二条路的最强形态——它不是"抓 AI 代写"（那是"查物质"，注定失败），而是：**看这个人自己的纵向曲线有没有出现无法由学习解释的断裂**。一个学生四个月前还写不出通顺的段落，今天交出结构精密的论证，中间没有任何可解释的过程——那是护照式的异常，不需要检出任何工具痕迹。

它把什么当成单独够用的那一样：路径。

四 · 第三条路：条件场裁定有效性（抗菌药物管理的读法）

青霉素今天的分子结构，和 1943 年一个原子都不差。

而在世界的很多地方，它已经不是那个药了。

这不是一个比喻，是抗菌药物管理这门学问每天在处理的事实。一种抗生素的临床有效性，不由分子决定，不由生产工艺决定，由**它周围那片微生物群落的敏感性谱**决定。而那片谱是被使用改变的。

这门学问最锋利的一条，也是它在本文里承担的那个位置，是这一句：**消耗有效性的不是滥用，是使用**。每一次完全正当的、指征明确的、剂量正确的处方，也在向那片场施加选择压力。滥用消耗得快一些，正当使用消耗得慢一些，但方向相同。所以这门学问的处方不是"别乱用"，而是"**把它当作一种会被用光的共有资源来管理**"——这一点在伦理学文献里被明确写成了公地悲剧的一个实例，并据此讨论过对抗生素使用征税。

第二条：成本被外部化到未来。 一张处方的记录里，有诊断、有药名、有剂量、有疗程、有转归。**没有一栏记"这一次用掉了多少还能再用的次数"**。那一笔落在未来的所有人身上，不落在任何一张处方上。

搬到毕业生身上：一个技能的价值不在这个人身上，在他所处的那片"还有多少人/多少机器能做同样事"的场里。而这片场正在被每一次使用改变。

它把什么当成单独够用的那一样：条件场。

五 · 三条路共同踩着的那块地

把三家的争论写成同一个句式，它是这样的：

三家争的是一份产出所显示的能力该归给谁、算多少。

化学说：归给这份产出本身，它显露成什么样就算什么样。反兴奋剂说：归给这个人，但要看过他的历史才能确认是不是他的。抗生素管理说：这个"算多少"由外部的场决定，与归给谁无关。

三家吵得很凶。但把下面这句话念给三家听，三家都会说"这还用说吗"：

能力是那种可以有归属的东西。

也就是说：存在一样东西叫"这个人的能力"，它长在某个地方（在产出里、在这个人身上、或在市场里），因而"它是不是他的""它值多少"这两个问题有答案，剩下的只是用什么方法把答案取出来。

三家分歧的全部内容，是用什么方法取。三家共享的是：**有一样东西在那里等着被取。**

这条前提要被推翻，才有别的东西能长出来。而按本文的纪律，推翻它的材料必须来自三家之一自己已经掌握、已经发表、只是没往这个方向用的事实——不能从外面搬一个第四家的立场进来。从外面搬进来的，只是加入了争论，没有取消争论。

推翻材料在第一家手里。

六 · 推翻它的那件材料：利托那韦，1998

本文的答案取自结构化学一侧，而被推翻的那条共有前提，也由结构化学自己的实证材料推翻。这个双重身份在这里写死，以免读者把它读成换皮。所用的材料是一组观察结果，不是化学家的立场——立场是要被推翻的那个东西，不能同时当推翻的工具。

1996 年，雅培把利托那韦作为治疗艾滋病的蛋白酶抑制剂投放市场，商品名 Norvir，半固体胶囊剂型。剂型这样设计，是因为开发期间只发现了一种晶型（Form I），而它在固态下口服吸收不好。

1998 年夏天，若干批次开始溶出度不合格。检查发现胶囊里析出了一种此前从未见过的晶型——Form II。它比 Form I 稳定得多，溶解度不到一半。药物随后被撤市，重新处方，损失以亿美元计。

这件事在结构化学里被反复讲，但通常讲的是"多晶型筛查要做全"。本文要用的是它的另外两处，两处都是纯粹的观察记录：

第一处：Form I 与 Form II 的分子式相同，连接方式相同，立体构型相同。**它们通过了全部的分子层面身份鉴定。** 差别在固态构象与堆积方式——也就是说，差别不在这个分子里，在这个分子是在什么条件下被析出来的。

第二处，这一处是刀：Form II 一旦在雅培的任何一间实验室或车间里成核过，**Form I 在那里就再也做不出来了。**不是难做，是做不出来。反复清洁不解决问题。这类现象在固态化学里有一个专门的名字：**消失的多晶型**（disappearing polymorphs）。它的定义就是：一种更稳定的新晶型一旦被发现，原有的那一种就变得几乎无法再制备。

把这两处合起来读：

"它是哪一样东西"不由它自己决定，由"它是在什么已经被别人长过的场里长出来的"决定；而这个场一旦被改变过一次，旧的那一样就不再是一个可达的选项。

这不是化学家的哲学主张，这是雅培的车间日志。

它推翻的正是第五节那条前提。因为它说的是：**不存在一样叫"这个东西本身"的、可以被归属的实体，等着人用某种方法把它取出来。**那个东西是什么，取决于它落地时那片场里已经有过什么。而那片场是有历史的，且历史不可逆。

三家为什么各自到不了这一步——注意，这不是三家的缺陷，是它们各自的问题结构使这一步成为不必要的：

- **结构化学**的问题是"这瓶东西是不是那样东西"，它需要一个与产地无关的判据，否则化学就不成其为一门可复现的学问。晶种场是它必须当作扰动来控制的東西，不是它要研究的对象。
- **反兴奋剂**的问题是"这个人该不该被取消资格"，它必须把判定锚在个体身上，否则处罚无法执行。"当时有多少人能跑出这个成绩"对它是完全无关的量。
- **抗菌药物管理**的问题是"下一个病人还有没有药可用"，它关心的是场的总体状态，个体处方只是加总项。"这一张处方的判别力"在它的问题结构里不存在。

三家的账本上，各自有一样东西被当成噪声、余数、或者根本不设字段：

家	不设字段的那一样
结构化学	这一批晶体是在什么已有晶种场里长出来的（晶种污染按扰动处理）
反兴奋剂	这个成绩被跑出的那一刻，世界上有多少人已能跑出同样的成绩
抗菌药物管理	这一张正当处方用掉了多少"还能再用的次数"

三条指向同一个空位：**每一次成功的执行里，"这一次让判别力少了多少"这一笔，没有字段。**

七 · 前提被推翻之后，显出来的是什么

AI 时代大学毕业生的核心竞争力，不是他掌握的技能存量，也不是他能力形成路径的可核性，也不是他相对于同侪的稀缺位置，而是他的一次产出中"这个世界还拿不出一个可复现的同类样本"的那一部分。

给这两部分各起一个名字，因为后面所有的引用都要靠它们才不会滑回旧概念：

未见份 (unseen share)：一次产出中，此刻的世界无法从公开可得资源里拿出功能等价样本的那一部分。

零判份 (null-discriminating share)：其余的那一部分。它在每一本账上都被记为这个人的贡献，而它在判别意义上是零。

三句话把零判份钉死，防止三种最容易的误读：

1. **零判份不是低质量的那一份。**它可能是全篇最漂亮的一段。质量与判别力是两根轴。
1. **零判份不是抄来的那一一份。**它可以是这个人完全独立、从头做出来的。他不知道世界上已有同类样本，这不改变判别力为零这个事实——就像雅培的技术员并不想做出 Form II。
1. **零判份不是"不重要的那一份"。**一份合格的行业报告，它的零判份可能承担了 95% 的实用价值。**判别力为零不等于使用价值为零。**这两件事被现行账本混成了一件，而它们的分离正是本文的全部内容。

为什么这是一样新东西，不是旧东西的新算法：删掉本文这个读数之后，零判份不是变得难测，是**不存在**。现行的能力账本——学分、绩点、作品集、简历条目、KPI、职级——全部把"产出"记为一个不分层的总量。零判份不是这个总量里的一个子集被记漏了；是这个

总量在设计上**没有可以把它分成两份的那一维**。学分不问"这门课学的东西世界上还有多少人会"，简历不问"你做的这件事此刻有几台机器能做"，作品集不问"这幅图的哪一部分是这个世界第一次看到"。三家的账各自都能记平，而这一笔谁也不欠。

八 · 辨别装置：一张二维格

只有一维就还是个名字。第二根轴是这样加的：

	未见份高	未见份低
通过形式审查	真竞争力（罕见；这一格里的人不需要解释自己）	靶格：合格的零判产出
未通过形式审查	未完成的原创（可救：补工艺即可）	普通的不合格

靶格是右上那一格：形式审查全部通过，而未见份接近零。

靶格必须是"全部形式审查都能通过"的那一格，否则它不需要一个框架——任何人都看得见。这一格的三条签名，逐条对照：

签名一 · 短期指标表现为改善。 落在这一格的学生，绩点更高、作品更多、交付更快、页面更整洁。生成式工具普及之后，这一格的所有可见指标全部向好。没有一个现行指标会因为一个人落进这一格而变差。

签名二 · 失效不产生任何信号。 这一格的人照常毕业、照常拿到面试、照常被拒；而拒信写的是"竞争激烈""与岗位匹配度不足"。没有任何一个环节会输出"你的未见份是零"这条信息，因为没有任何一个环节在测它。招聘方自己也说不清为什么这个候选人"看着都对，就是不想要"。

签名三 · 自我加固：越正规化越严重。 这一条最要命。课程标准越明确、评分细则越可量化、作品集要求越具体、能力框架写得越清楚——产出就越向"已有可复现样本"的那一侧集中。因为一条被写清楚的标准，本身就是一份可复现的规格说明；照着它做出来的东西，按定义落在世界已经能拿出样本的那一侧。**把培养目标写得更清楚，是把学生更快地推进靶格。**

三条签名齐，靶格成立。

九 · 读数：未见率 ρ

未见率 ρ = 未见份 / 一次产出的全部，取值 $[0, 1]$ 。

这是一个无量纲的比率，因此它可以在量纲完全不同的三个领域里并排比较——这也是本文对它作为读数的第一条要求：

领域	一次"产出"是	未见份怎么数	量纲
结构化学	一次合成	该分子的这一晶型此前有无被报告过的可复现制备	晶型条目数
竞技体育	一次成绩	该成绩在当年世界排名中位于何处，及此前有多少人达到过	人次
抗菌治疗	一次处方	该药对该病原在该地区的现行敏感率	百分比
毕业生的一份报告	一份交付	其中有多少段落无法由公开资源在合理时间内产出功能等价物	段落数或工时

四处量纲不同，比率相同。

第四条性质是最容易漏也最要命的： ρ 必须与既有主要指标正交。 判据是：举得出一个“既有指标最好看、而 ρ 最难看”的真实例子。举得出——一份拿了满分的课程作业，可以有 $\rho \approx 0$ ；反过来，一份被判不及格的作业，可以有很高的 ρ （它做的事没人做过，因而没有评分细则能接住它）。举不出这样的例子，这个读数就该重造。这里举得出，而且这两种情况在任何一所大学里都随处可见。

分层引用结构（本文可以被分层引用，这是诚实，也是保护）：

- **第一层**（最强、最易倒）：未见份是 AI 时代竞争力的承重判据。
- **第二层**（分析工具）：上面那张二维格，及其靶格三签名。
- **第三层**（最稳、最容易检验，且不依赖前两层）：**判别力的衰减速率主要由暴露时长决定，与使用强度关系微弱，且不因保密而停止。** 第十一节检验的就是这一层。

九之二 · 是哪一步把零判份压掉的，又是什么条件使这一步不被察觉

前一节交出的是一个空位：一次产出里有一份东西，在所有账本上都找不到它的位置。**只交出空位是不够的。** 一个空位若不能说清它是被哪一步操作压掉的、以及什么条件使这种压掉可复现，那它就只是一个命名。

本文的承重命题剥到形式层是这样一句：**显露上缺了一块，是路径上的某一步与条件场里的某片条件合起来把它压掉的。** 这一句与一族既有工作同形，本文与那一族的分界写在第十二节之二。

甲 · 空位

零判份。三家的账本里为什么找不到它，见第六节末尾那张表。

乙 · 压掉它的那一步操作，以及它为什么是对的

那一步是**总量归集**：评价系统在收下一次交付时，把它的全部承载单元合并成一个分数、一个学分、一条简历行、一个交付记录。

这一步在它自己的范围里是正确的，而且不是勉强正确，是理由充分地正确：

1. **评价必须可比、可加总、可申诉。** 一个分数能排序，一组分层判定不能。
1. **分层记账贵一个数量级。** 按第十节的手册，每一个承载单元要两名不知作者身份的编码者独立判定加一个仲裁。一份作业的评分成本从十分钟变成两小时。
1. **最要紧的一条：在复制成本高的世界里，总量归集是无损的。** 那时几乎每一个承载单元都落在未见份里，分层与不分层给出**同一个排序**。一个曾经无损的简化，不需要理由就会被保留下来。

检验句（这一条决定它是机制还是事故）：把这一步改对——改成分层记账——空位会不会消失？**会消失，但代价是评价成本上升一个数量级，并引入编码者分歧这一新的不公平来源。** 会因为把它改对而消失、且改对的代价是结构性的，这说明它是机制，不是失误、不是疏忽、也不是资源不足。**没有任何一位评分者做错了什么。**

丙 · 使这种压掉不被察觉且可复现的条件场

三条，缺一条这种压掉就会被发现：

1. **复制成本的下降是连续的、无事件的。** 没有哪一天有人宣布"从今天起这一类产出不再稀有"。一条没有台阶的曲线不产生可被记录的时刻，因而没有任何一个环节需要为它做出反应。
1. **评价与兑现之间隔着时滞。** 学分在毕业时结算，判别力在就业市场结算，中间隔两到四年。时滞使反馈无法回到评价环节——等到市场给出信号时，给出这个分数的那门课已经又开过三轮。
1. **零判份的使用价值仍然为正。** 这一条最容易被漏掉，也最关键：如果零判份是无用的，学生、教师、雇主三方都会立刻察觉。可它是有用的——那份报告确实解决了问题。**所有当事人的短期体验都完全正常**，这正是第八节靶格签名二"失效不产生任何信号"的机制来源。

换一片条件场，它还会被压掉吗？ 不会——而这正是本文可被证伪的一个入口。在复制成本高且不下降的领域，三条条件不成立：需要受限设备的实验工作、以受限数据为唯一输入的分析、必须在场完成的照护与服务。这些领域的总量归集至今仍是无损的，零判份在那里接近零，因而它不作为一个空位存在。**这与第十六节写的"本文不适用的边界"是同一条线，从两头各走一遍。**

十·清点手册

一个读数如果粒度不写死，它就不是读数。以下四项在本文正文、证伪条款、赌注三处引用同一个定义、同一个阈值。

读数名：未见率 符号： ρ

定义： $\rho = \text{未见份} / \text{产出总量}$ 。

产出总量 = 该次交付中可独立计数的最小承载单元总数。

粒度：一个"承载单元" = 一个能独立完成一件事的段落、函数、图表或设计决策。边界例：一段三句话的方法说明，若删去它读者无法复现该步骤，则算一个单元；若它只是复述上一段，不算。

编码规则（五条）：

1. 判定基准时点：交付日。此后世界的变化不回溯修改该次 ρ 。
2. "可复现同类样本" = 在交付日，一个具备该领域本科水平、可访问公开互联网与公开可用生成式模型的人，在 30 分钟内能产出功能等价物。30 分钟是本文全篇唯一阈值。
3. 功能等价 = 在该单元所服务的那个判断上给出相同结论，且论据强度不低于原件。措辞不同不影响判定。
4. 判定由两名不知作者身份的编码者独立进行，分歧交第三人。
5. 不可判定的单元单列一栏，不计入分子也不计入分母，并报告其占比。占比 $> 20\%$ 时该次 ρ 作废。

表头：| 单元号 | 内容摘要 | 编码者A | 编码者B | 仲裁 | 计入未见份 Y/N |

不适用：不足 10 个承载单元的交付；纯执行性任务（数据录入、格式转换）；以及以保密数据为唯一输入的交付（其未见份来自数据准入而非产出本身）。

三处引用一致性自查：正文 § 9-10 = 证伪 § 13 = 赌注 § 15

十一 · 一次真实检验，及它的不利结果

证伪条款不许全是纸面的。本文用公开数据跑了最便宜的那一条——第三层主张：**判别力的衰减主要由暴露时长驱动，而不由使用强度驱动。**

预注册（在读取检验数据之前写死，此处照原样抄录）

检验对象：Ott 等 2022 年发表于 *Nature Communications* 的 AI 基准饱和数据集。

H1：基准从发布到饱和的时间尺度以年计，不以十年计。

H2（本文最愿意押、也最容易翻车的一条）：使用强度不是主要驱动量——被用得更多的基准，并不显著更快饱和。若该文报告“使用量与饱和速度显著正相关”，本文第三层主张被驳，须退回。

H3：饱和基准占比落在 10%-90% 之间。近 100% 则不是机制而是必然律；低于 10% 则消耗不普遍。

停止规则：只读该文摘要与结果段，读一次，不换数据源，不追加检验。

不利结果处置：无论方向如何照原样写进正文，并写明它改了本文的什么。

结果

该文覆盖 3765 个基准，横跨计算机视觉与自然语言处理两个完整领域。它报告：**很大一部分基准迅速趋向接近饱和；许多基准从未获得广泛使用；不同任务的性能增益容易出现难以预见的突进。**

逐条对照：

- **H1：获定性支持。** 该文明确使用“迅速”（quickly）描述趋向饱和的过程，其数据覆盖的时间窗以年计。
- **H2：本次真跑中未获裁决。** 该文的公开摘要与结论段并不包含“使用强度 × 饱和速度”的直接检验。按预注册的停止规则，本文不追加数据源、不改换主判决量。**H2 在本次检验中没有结果。**
- **H3：未获定量裁决。** 该文给出的是定性表述（“很大一部分”），没有可直接读出的百分比。

这次真跑的信息量低于预期，这是一个不利结果。 本文照原样报告它，并写明它改了什么：

改动一：第三层主张从“已获检验支持”降级为“**有一处事后观察，尚未经预注册检验**”。那处事后观察是：一项 2026 年对 60 个语言模型基准的系统研究报告，在控制基准年龄之后，引用数（ $\rho=0.22$, $p=0.12$ ）、引用增长率（ $\rho=0.13$, $p=0.37$ ）与在厂商技术报告中出现的频次（ $\rho=0.05$, $p=0.73$ ）**均与饱和无显著关联**；而私有测试集（ $N=4$ ）与公开测试集（ $N=56$ ）的饱和分布**没有可辨的差异**。这两条都指向本文的方向。**但它们是我在写下预**

注册之前就已读到的，因此它们不是检验，是线索。按本文自己的纪律，线索不能充当检验。

改动二： ρ 的可测性评级下调。原本预期存在现成的公开数据可以直接读出"衰减速率"，现在确认没有——**这个量必须自建清点**。清点手册（第十节）因此从"辅助材料"提升为本文的承重件之一：一个必须自建清点的读数，其成立与否首先取决于清点规则写得够不够死。

改动三：本文承认一处它解释不了的洞。上述 2026 年研究同时报告，**测试集规模越大，饱和指数越低**。这一条无法归入"消耗"，它是测量分辨率问题——分辨率不足会伪装成饱和。本文目前没有办法在一次交付的 ρ 上把这两者分开。这是本文最脆弱的一处。

十二 · 与最近的几种说法的分界

七条，每条给出判定形式。判定形式的格式是死的：**若观察到 X，则本判断对而该说法错；若观察到 Y，则反之。**

（一）信号理论（Spence 1973，经济学）

它说到哪一步：教育是一个成本与能力负相关的信号，均衡中高能力者选择更高的教育水平，从而与低能力者分离。

分离线：信号理论关心的是**分离均衡是否存在**；它假定信号的成本结构是外生给定的。本文说的是**信号的判别力被时间消耗**，而消耗与谁在发信号、发了多少无关。

判定形式：若观察到一个**从未有人使用过**的新信号，在无人使用的情况下，其判别力随时间下降，则本判断对而信号理论错；若判别力只在被大量使用之后才下降，则反之。

（二）位置性商品（Hirsch 1976，经济学，专著层）

它说到哪一步：某些物品的价值来自其相对位置，其稀缺无法靠技术进步缓解；高等教育是主要例子。

分离线：位置性商品说的是**分配问题**——不是所有人都能进哈佛。本文说的是**耗竭问题**——一个人的产出的判别力，会被"世界上出现了可复现样本"消耗掉，即使没有任何人与他竞争同一个位置。

判定形式：若观察到一个**没有竞争者**的场景（唯一候选人、无名额上限），其产出的判别力仍随时间下降，则本判断对而位置性商品说错；若无竞争时判别力不变，则反之。

（三）文凭社会 / 文凭通胀（Collins 1979，社会学，专著层）

它说到哪一步：随着持有某文凭的人数增加，该文凭的市场回报下降，导致教育要求不断攀升。

分离线：文凭通胀是**存量稀释**——分母变大。本文说的不是分母变大，是**分子里有一部分本来就是零**，而这部分与持有人数无关。一个全世界只有他一个人拿的学位，其对应的产出照样可以 $\rho \approx 0$ 。

判定形式：若观察到一个**持有人数为一**的资格，其对应产出的判别力仍为零，则本判断对；若判别力只随持有人数上升而下降，则反之。

（四）古德哈特定律 / 坎贝尔定律（Goodhart 1975 / Campbell 1976）

它说到哪一步：一个指标一旦成为目标，就不再是好指标；被优化的测量会失效。

分离线：这是**被博弈**导致的失效——机制是有人朝着指标使劲。本文说的失效机制是**正当使用**：没有任何人博弈，每个人都诚实地做，判别力照样被消耗。这与第四节抗菌药物管理那一条同构——消耗有效性的是使用，不是滥用。

判定形式：若观察到一个**无人有动机去优化**的判据（例如一个没有任何后果的匿名测评），其判别力仍随时间下降，则本判断对而古德哈特错；若只有在有优化动机时才下降，则反之。

（五）题目曝光控制与题目退役（Simpson & Hetter 1985；Stocking & Lewis 1995；心理测量学·方法学占位者）

它说到哪一步：这是本文最危险的一位。自适应测验里，判别力高的题目被选中的频率最高，因而最快被过度曝光、被泄露、被记住，必须限制曝光率甚至退役——这门方法学明确地把“判别力会被使用消耗”写成了工程问题，并给出了六种以上的控制算法。

分离线：两处，都可判定。**第一处·机制不同**。题目曝光控制假定消耗机制是**考生的不当获取**（记住了、传出去了），因此它的对策是保密、随机化、限制曝光率。本文说的机制是**世界上出现了可复现样本**，与获取无关；因此保密不起作用。**第二处·对策不同**。题目曝光控制的处置是**换题**，它假定题库外还有干净的题。本文的推论是换题救不了——新题一旦被执行一次，同样开始消耗。

判定形式：若观察到**私有的、从未泄露**的判据仍以与公开判据相同的速度失去判别力，则本判断对而曝光控制的机制假设错；若私有判据显著更慢，则反之。**关于这一条，第十一节引的那项 2026 年研究给出了一个方向一致的线索（私有 N=4 与公开 N=56 的饱和分布无可辨差异），但样本量太小，且它不在本文的预注册里，因此本文不据以宣称此条已经判定。**

（六）撒谎者红利（Chesney & Citron 2019，法学）

它说到哪一步：一旦深度伪造广泛存在，说谎者可以把真实的证据斥为伪造，从而逃避追责；伪造品的存在损害了真品的可信度。

分离线：这是本文形状上最近的一位，但方向相反。撒谎者红利说的是**真品的可信度被伪品的存在拉低**——损害的是真伪判定。本文说的是**真品的判别力被同类真品的存在耗尽**——一份完全真实、完全由人写的报告，其判别力照样归零，与真伪无关。

判定形式：取一份**其真实性已被完全确证**（有完整过程留痕、有第三方见证）的产出。若它的判别力仍然为零，则本判断对而撒谎者红利错；若确证真实性之后判别力恢复，则反之。

（七）个体参照评价（Hughes 2011, 2014，教育评价·应用领域实务正主）

它说到哪一步：把评价的参照点从“与他人比”和“与标准比”改为“与自己的上一次比”，据此评价进步而非成就。这是教育评价界对本文所述困境给出的最成熟的现成处方，且它与第三节生物护照的转向是同一个动作。

分离线：个体参照评价把锚点挪到了**这个人自己的历史**，因而它测的是**成长**。本文说的量锚在**这个人与世界现存样本之间**，因而它测的是**差**。一个人可以在个体参照下持续进步（每一次都比上一次好），而他的 p 一路走向零——因为世界跑得更快。这两个量可以同时是真的，且方向相反。

判定形式：取一组学生，同时测个体参照的进步幅度与 ρ 。若观察到进步幅度为正而 ρ 单调下降的个体占相当比例，则两者确为不同的量，本判断成立；若两者高度相关，则本文的量是它的一个代理，应当合并。

最近的邻居是第五位——题目曝光控制。 本判断仍不可被它吸收，理由是可判定的那一条：它的整套对策建立在“保密能延缓消耗”之上，而本文的推论是保密不起作用。这两条不能同时为真。

划完之后再取一次交集：上面七位各自结构性地看不见什么？

- 信号理论把“信号的判别力”当作给定，因此它看不见“信号自己在衰减”；一旦承认，均衡就不再是均衡，而是一个必须持续再生的过程。
- 位置性商品把“位置的总数”当作给定，因此它看不见“位置本身会消失”。
- 文凭通胀把“资格的所指”当作给定，因此它看不见“同一张文凭前后指的不是同一件事”。
- 古德哈特把“优化动机”当作给定，因此它看不见“无动机的正当使用也在消耗”。
- 曝光控制把“题库外还有干净的题”当作给定，因此它看不见“新题一执行就开始消耗”。
- 撒谎者红利把“真伪是判别的核心”当作给定，因此它看不见“确证为真也救不回判别力”。
- 个体参照评价把“与自己比是安全的”当作给定，因此它看不见“自己进步而差在扩大”。

这七行背后是不是同一个“被当作给定的东西”？是。它是：

一个判据的有效性，是它被使用时的一个属性；不使用它，它就原样保存在那里。

所有七位——包括互相批判的各派——都站在这一条上。这就是为什么“没有人辨识过”：**不是没人看，是看的人都站在它上面。**

十二之二 · 与同族、与同批的分界

第十二节处理的是学术史上的近邻。还有两类更近的东西必须处理，而它们通常被略过：**与本文同属一族的既有工作，以及与本文同时出自同一条产线的兄弟篇。**这两类是本文最近的占位者，略过它们等于把最容易的一条反驳留给读者。

一 · 与同族：「账上没有那一栏」

本文的形式句与一族数量可观的既有工作同形：某样东西在正式记录里没有字段，因此它不被计入，因此它在制度上不存在。这一族的处方是一致的——**加一栏。**

分界：那一族说的是**字段缺失**；本文说的是**字段齐全、记的量也正确，而那个量在设计上不可分层**。学分记录并没有漏记什么，它记的是“这个人交了什么、有多好”，这两件事都记对了。缺的不是一栏，是一根轴。

判定形式：在一本能力账本上新增字段——例如“AI 使用披露”栏、“过程留痕”栏、“工具清单”栏。若观察到新增字段之后判别力恢复，则同族对而本文错；若观察到新增之后判别力照旧按暴露时长下降，则本文对。**这一条已经在被现实检验：过去两年里，“AI 使用披露”栏在多少所高校被加上了。它加的是第二条路（看路径），不是本文的轴。** 本文预测它不会使判别力恢复。

二 · 与同批：同一道题的另外两次落地

本文并不是这道题唯一的一次回答。同一时期，同一道题另有两次独立落地，三次的三个源互不重叠，得到的也是三个不同的东西。三者不是三种说法，是三个测量发起处：

	读数锚在哪两样之间	读数
现配优势	人与场之间——优势是每次交手当场由构件加配组做出来的，不是随身携带的存量	留种率 s ：完备交接清单被同资历者执行后，复现原表现的份额
零判份（本文）	产出与世界现存样本之间	未见率 p
脱稿步	产物与走出它的那条路之间——档案在，而走出它的那一步没有记录	脱稿成立率 q ：已发起的追问中，当场走出可核下一步的份额

与现配优势的判决性对照预测：取一个人在同组织、同配组、条件完全不变下连续五年的交付。现配优势预测显露表现稳定——配组没变，优势就还在。本文预测未见率单调下降——世界在变，配组不变救不了。若观察到表现稳定而判别力仍在下降，本文对；若两者同步变化，本文被现配优势吸收。

与脱稿步的判决性对照预测：这一条的分界不在结论上，在测量发起处。本文读的是产出本身，脱稿步读的是一次被发起的追问。因此上一段那个“五年不变”的情形，脱稿步对它不作预测——没有人发起追问时它没有读数。反过来：取一份未见率极低的交付，对作者发起一次锚定其自身产物的追问。若脱稿成立率高而未见率低，两者确为不同的量；若两者高度相关，本文的量是脱稿步的一个代理，应当合并。

三者若在同一份台账上给出相反评价，取发起过追问的那一份为准——因为那一份多了一次真实的观察，而另外两份只读了留下来的东西。

互补而非竞争的那一处（这比“我们不一样”值钱）：现配优势解释为什么带走一个人不等于带走他的表现；本文解释为什么留住一个人也不等于留住他的判别力。两者的分子分母恰好错开：前者问的是优势能不能被搬运，后者问的是优势会不会自己蒸发。一个组织可以同时在这两件事上失手，而它们的叠加不是相加——留种率低使表现不可迁移，未见率低使表现不可辨识，两者同时发生时，“这个人到底行不行”这个问题连一个错误答案都得不到。

十三 · 证伪条件

最强的竞争解释先写出来，不挑软柿子：*其实不就是“会的人多了就不值钱”吗？*

这是最朴素、最难反驳的那一句。它由供给增加解释一切，不需要本文任何机制。要排除它，必须拎出一个只有本文机制成立才会出现、而供给解释预测不出的变量：**衰减在供给不变时是否仍然发生。**

第一条（主检验，格式写死）

控制“具备同等技能的人数”与“该领域从业者总数”之后，若“该产出类型进入公开可复现语料的时长”与其判别力之间不存在剂量-反应关系，则本判断为伪——届时困境应归因于人员供给而非样本可复现性。

第二条（保密检验，这是与第五位邻居的分水岭）

取同一时期、同一领域、判别力初始值相当的两组判据，一组公开、一组严格保密从未外泄。若保密组的判别力衰减显著慢于公开组，则本文为伪。

第三条（无竞争检验，与第二位邻居的分水岭）

取一批不存在名额上限、不存在竞争者的评价场景（例如资格认定而非选拔）。若其中产出的判别力不随时间衰减，则本文为伪。

第四条（正向读数，用"顺序早于"而不是相关）

在同一机构内，若"某类交付的 ρ 中位数开始下降"这一事件 **不早于**"该类岗位的入门招聘量开始下降"这一事件，则本文的因果方向为伪。相关关系容易被巧合满足，顺序难得多。

第五条（低成本可实做，且所需数据在多数机构的既有记录里已经存在）

任何一所大学的教务系统里都存有近五年同一门课程的作业文本。取同一题目、同一评分细则下 2021 年与 2026 年的两批作业，按第十节的手册各计一次 ρ 。若两批的 ρ 分布无可辨差异，则本文的核心经验主张在教育场景中为伪。 **这一条不需要新数据、不需要伦理审批之外的任何资源，一个研究生两周可以做完。**

样本纪律：以上凡涉及跨案例比较者，样本须 $N \geq 6$ 且在关键混淆变量上同质。 **严禁拿两个国家或两个地区充数**——那类 $N=2$ 的异质对比在方法上站不住，只能作启发式说明。优先在同一制度环境内找自变量有真实变异的多个单元：同校的多个院系、同行业的多家企业、同一平台的多个社群。

这条判断若被证伪，我们会学到什么：如果第二条（保密检验）翻车——保密确实能显著延缓衰减——那么判别力的消耗机制就是获取而不是可复现性，处方随之完全不同：应当去建判据的保密基础设施，而不是去建"持续产生未见份"的能力。那是一个方向相反、但同样可执行的结论。

十四 · 本文对自己的定位

不是说"本文也有局限"。是用本文自己的判据给本文自己定位，并说出这个定位的具体后果。

本文自己就在那张格子里，而且它必须诚实地报告自己落在哪一格。

本文的承重命题是一个判据。按本文自己的说法，一个判据一旦进入可及状态，它的判别力就开始按暴露时长衰减，且保密不起作用。 **因此本文在发表的那一刻就开始失效。**这不是修辞。具体后果有三条：

后果一：若本文有价值，则它的价值在被引用得最多的时候已经最低。一条被大量引用的判据，正是暴露时长最长、被最多写作者内化过的判据；到那时，"用未见份来看人"这句话本身会成为可复现样本，说出它不再证明说话者看见了什么。

后果二：本文的三个源——结构化学、反兴奋剂、抗菌药物管理——之所以还能长出新东西，恰恰因为它们对"毕业生竞争力"这个问题的三句话尚未被就业能力研究收编。 **本文写完并公开之后，这个条件被本文自己破坏了一部分。**下一个想用这三家来答这道题的人，未见份比我低。这与第六节的 Form II 是同一个形状：一旦在这间厂房里长过一次，旧的那一种就不再是可达的选项。

后果三，也是最不舒服的一条：本文是用一套工序做出来的。工序意味着可复现。因此本文自己的 ρ ，很可能低于它读起来的样子。按第十节的手册严格计，本文的承载单元里有多少能在三十分钟内被一个具备本科水平、可访问公开资源的人做出功能等价物？我不知道，因为我不能对自己的稿子做双盲编码——本文自己的纪律不允许作者给自己的文本发认证分。这个数必须由别人来数，而在它被数出来之前，本文关于自己的一切说法都只是设计目标。

十五 · 一条写死日期的赌注

赌注：到 2029 年 9 月 1 日为止，至少有一所设有本科教育的高校，在其正式的、对学生公开的评价文件里，出现一栏与本文第十节手册结构等价的字段——即：**要求就一次交付分别记录“世界已能复现的部分”与“世界尚不能复现的部分”，并把后者单独计分。**

什么不算命中（这一条与赌注本身同等重要）：

- 校长讲话、白皮书、战略规划、研究报告里出现类似说法 —— **不算。**
- 某位教师在自己课上的个人做法 —— **不算。**
- 评分细则里出现“原创性”“创新性”“独特性”这类词 —— **不算。**这些词在多数评分细则里已经存在几十年，它们评的是质量，不是判别力；两者是第八节那张格子的两根不同的轴。
- 只把“是否使用 AI”作为一栏 —— **不算。**那是过程审查，是第二条路，不是本文。
- 必须是**可执行的评价条款**：有编码规则、有计分方式、有申诉通道。**只要求机构自行说明的，不算命中。**

我押“会命中”。我同时写下我认为最可能的失败方式：这个字段太贵。它要求两名不知作者身份的编码者独立判定每一个承载单元，成本远高于给一份作业打一个分数。**贵是它最可能死掉的原因，而不是错。**

十六 · 不适用的边界，与伦理约束

本文不适用于：

1. **纯执行性岗位与纯执行性任务。** 数据录入、格式转换、按既定规程操作——这些工作的价值从来不在判别力上，本文对它们没有任何说法，把 ρ 用在它们身上是误用。
1. **以受限数据为唯一输入的交付。** 一份基于保密数据的分析，其未见份来自数据准入，不来自产出本身。本文的量在这里读不出有用的东西。
1. **服务性、照护性、在场性的劳动。** 一次护理、一次安抚、一次课堂上的临场判断——它们的价值在于发生在那个人和那个人之间，本身就不以“可否被复现”计价。本文的整个框架在这一类劳动面前是失效的，而这一类劳动占人类劳动的很大一部分。
1. **未成年人的学习过程。** 见下。

反过来被加的约束（两处，都真削弱了本文原本想主张的东西）：

第一处 · 来自抗菌药物管理。 那一家坚持一件事：**有效性是共有资源，个体不为它的消耗负责。** 一张正当的处方消耗了未来的有效性，但没有人会因此指责开处方的医生。这一条直接毁掉了本文最自然的一个处方——“教学生去追求高未见份”。如果没有这一处约束，我原本会写下的那句更强的话是：“**学生应当为自己产出的 ρ 负责。**” 这句话现在不能写了。 ρ 的下降主要由世界的变化造成，把它记在个体头上，与因为耐药率上升而指责病人是

同一种错。本文的适用范围因此从"评价个人"缩小到"诊断系统"。这是一处价值排序的改变，不是一个补充条件。

第二处 · 来自反兴奋剂。那一家用六十年学会了一件事：一套用于取消资格的判据，一旦存在，就会被用于取消资格。生物护照的每一次误判都是一个人的职业生涯。这一条迫使本文放弃另一个原本很自然的主张——把 ρ 用作**准入筛选**。它不能这样用，理由不是效果不好，是**代价的分布**：筛掉一个 ρ 低的应届生，他失去的是入场机会，而他的 ρ 低主要不是他的过错（见上一处）。因此 ρ 只能用于**诊断与设计**，不能用于**选拔与淘汰**。

伦理三条（本文的产物是一个比率，而比率几乎必然会被拿去考核人，所以它自带伦理件）：

1. **这个读数指的是位置，不是人。** ρ 描述的是"一次交付与世界现存样本之间的差"，不是"这个人有多好"。把某人称为"低 ρ 的人"是对本读数的误用。
1. **不得为了把这个数做好看而制造无谓的差异。** 具体地说：不得鼓励学生刻意选择冷僻、无人问津、或故意异常的做法来抬高 ρ 。**冷门不等于未见份**——一个没人做过的蠢事， ρ 很高而价值为零。第八节那张格子有两根轴，只优化一根是对框架的背叛。
1. **已经按这个读数做出不利安排的，必须公开编码规则并给出复核通道。** 一个人有权知道自己的哪一个承载单元被判为零判份、依据是什么、由谁判的，并有权要求重判。

反噬预言（工序要的正是这一条）：本命题一旦被制度采用，最省事的实现方式恰恰会摧毁它要保护的那个东西。具体路径是可预测的：学校会把"提高学生的未见份"写进培养方案，于是需要一份**可执行的操作指南**——"如何做出世界还没有的东西"。而这份指南一旦写出来并被推广，照着它做出来的东西就落在世界已经能拿出样本的那一侧。一份关于**如何产生未见份的标准化指南**，是**零判份的高效生产线**。

我看不到出口。这不是谦辞，是本文的一个真实缺口：任何可传授的方法都是可复现的，而可复现的方法产不出未见份。看得到出口的反噬不是反噬，是待办事项；这一条我看不到出口，所以我把它原样留在这里。

结论

三条路各自说对了一部分：成品确实是唯一可测的东西，路径确实决定了它是怎么来的，条件场确实决定了它还算不算数。它们错的地方在同一处——**它们都以为有一样叫"这个人的能力"的东西在那里，等着被某种方法取出来。**

利托那韦的车间日志说的是另一回事：那样东西是什么，取决于它落地时那片场里已经有过什么；而那片场有历史，历史不可逆。

于是"AI 时代大学毕业生的核心竞争力是什么"这个问题，它真正的答案不是一份技能清单，也不是一套过程规范，甚至不是一个位置。它是一个**差**：这个人交出来的东西，与此刻的世界已经能拿出来的东西之间，还剩多少距离。

这个差每天都在收窄。它不因为你把它锁起来而停止收窄，也不因为没人使用而保持原样。它是这样一种东西：**你每一次成功地证明自己，都在花掉证明自己所用的那把尺子。**

而现在的每一本能力账本，都没有记这一笔的字段。

注释与文献

结构化学一侧

- Bauer, J., Spanton, S., Henry, R., Quick, J., Dziki, W., Porter, W., & Morris, J. (2001). Ritonavir: An extraordinary example of conformational polymorphism. **Pharmaceutical Research**, 18(6), 859–866.
- Chemburkar, S. R., et al. (2000). Dealing with the impact of ritonavir polymorphs on the late stages of bulk drug process development. **Organic Process Research & Development**, 4(5), 413–417.
- Morissette, S. L., Soukasene, S., Levinson, D., Cima, M. J., & Almarsson, Ö. (2003). Elucidation of crystal form diversity of the HIV protease inhibitor ritonavir by high-throughput crystallization. **PNAS**, 100(5), 2180–2184.
- Bučar, D.-K., Lancaster, R. W., & Bernstein, J. (2015). Disappearing polymorphs revisited. **Angewandte Chemie International Edition**, 54(24), 6972–6993.
- Parent, S. D., et al. (2023). Ritonavir Form III: A coincidental concurrent discovery. **Crystal Growth & Design**, 23(1), 320–325.

反兴奋剂科学一侧

- Sottas, P.-E., Robinson, N., Rabin, O., & Saugy, M. (2011). The Athlete Biological Passport. **Clinical Chemistry**, 57(7), 969–976.
- World Anti-Doping Agency. **Athlete Biological Passport Operating Guidelines** (current version).

抗菌药物管理一侧

- Giubilini, A. (2019). Antibiotic resistance as a tragedy of the commons: An ethical argument for a tax on antibiotic use in humans. **Bioethics**, 33(7), 776–784. DOI: 10.1111/bioe.12598
- Hardin, G. (1968). The tragedy of the commons. **Science**, 162(3859), 1243–1248.

近邻（第十二节逐条交手）

- Spence, M. (1973). Job market signaling. *The Quarterly Journal of Economics*, 87(3), 355–374.
- Hirsch, F. (1976). *Social Limits to Growth*. Harvard University Press.
- Collins, R. (1979). *The Credential Society: An Historical Sociology of Education and Stratification*. Academic Press.
- Goodhart, C. A. E. (1975). Problems of monetary management: The U.K. experience. Papers in Monetary Economics, Reserve Bank of Australia.
- Campbell, D. T. (1976). Assessing the impact of planned social change. Occasional Paper Series #8, The Public Affairs Center, Dartmouth College.
- Sympson, J. B., & Hetter, R. D. (1985). Controlling item-exposure rates in computerized adaptive testing. *Proceedings of the 27th Annual Meeting of the Military Testing Association*, 973–977.
- Stocking, M. L., & Lewis, C. (1995). *Controlling Item Exposure Conditional on Ability in Computerized Adaptive Testing* (ETS Research Report RR-95-24).
- Chesney, R., & Citron, D. K. (2019). Deep fakes: A looming challenge for privacy, democracy, and national security. *California Law Review*, 107, 1753–1820.
- Hughes, G. (2011). Towards a personal best: A case for introducing ipsative assessment in higher education. *Studies in Higher Education*, 36(3), 353–367.
- Hughes, G. (2014). *Ipsative Assessment: Motivation through Marking Progress*. Palgrave Macmillan.

背景与检验数据

- Polanyi, M. (1966). *The Tacit Dimension*. Doubleday.
- Autor, D. H. (2014). Polanyi's paradox and the shape of employment growth. NBER Working Paper No. 20485.
- Agrawal, A., Gans, J. S., & Goldfarb, A. (2019). Exploring the impact of artificial intelligence: Prediction versus judgment. *Information Economics and Policy*, 47, 1–6.
- Ott, S., Barbosa-Silva, A., Blagec, K., Brauner, J., & Samwald, M. (2022). Mapping global dynamics of benchmark creation and saturation in artificial intelligence. *Nature Communications*, 13(1), 6793. DOI: 10.1038/s41467-022-34591-0 — 本文第十一节的预注册检验对象。
- Akhtar, M., Reuel, A., et al. (2026). When AI benchmarks plateau: A systematic study of benchmark saturation. *Proceedings of the 43rd

版本与分工声明

字数：15,461 汉字（机检，含标题与文献表中的汉字），低于常规篇幅，此处如实标注，未作补白。

人机分工：三个领域的选择、共有前提的识别、推翻材料的定位、承重命题的形成、二维格与读数的设计，由本次写作过程完成；第六节的利托那韦事实、第十一节的两份基准研究数据、第十二节七位近邻的出处，均经公开检索核对，出处列于上；第十一节的预注册在读取检验数据之前写下，其未获裁决的结果照原样报告，未事后调整判负条款、未改换主判决量。

本文不给自己打分。 上文任何关于本文质量的表述都是设计目标，不是认证结果；须交独立评审盲评复核。

测量层交代了读数在哪里有定义，还剩最后一问：一份落在定义域内的读数，凭什么可信？《他手无事账》守着这一关，它的回答只有一根轴——**取样时刻归谁定**。

自己选好时刻交出的证据，在准备可以外包的时代同时失效；未准备的时刻是唯一还带信息的时刻。反兴奋剂、民航、消费信贷三个行业各花了一代人撞明白这件事，又各自只拿到答案的三分之二：三条归属——时刻归谁定、取样之手归谁、记录归谁持有——张成八格，三家与毕业生各占七格，有一格四方皆空：**他人之手、非本人所选时刻、无异常、且记入本人可携带的账本**。这一章的承重物就是那格空位，以及一个可以拿去问任何一份简历的判据："在这个人的能力证据里，有几份的取样时刻不是他自己定的？"

它与第四部的接口一句话说清：常席差告诉你**去哪里读**（跨配组），他手无事账告诉你**凭什么信你读到的**（他择时刻）。两轴正交——配组轮换但全部预约，与配组固定但无预告抽查，是两个独立变动的数。

第五部 取证：证据层

第十三章 他手无事账

【章首】本章给可读条件立第四条，也是取证层的独章：**他择时刻**。三个此前二三十年各自撞上“产物可被外力制造”的行业——反兴奋剂、民航、消费信贷——共同默认“存在一处不能被制造的参照物”，又各自被自己的数据推翻；推翻之后剩下的不是第四处参照物，而是三条从未被当作对象的归属：取样时刻归谁定、取样之手归谁、记录归谁持有。三条张成八格，三家与毕业生占了七格，有一格四方皆空：他人之手、非本人所选时刻、无异常、记入本人可携账本——本章以那格空位为承重物，读数是他择率与归己率。本章的三次预注册审计里，第二次的读后改判（承重物由一条属性改到那格空位）连同改判理由原样照登。本章欠的账：他择读数与入职绩效的剂量-反应机制检验——设计已备，最便宜的版本只需一家愿意回看自己招聘档案的机构；另欠医学教育工作场评估一族（受训者自选时机与观察者，恰落空格上方一行）一段点名分界，本书第十五章代为补记。

大白话：自己挑的日子请人来检查，不算数

一家餐馆自己请了美食评委，选了个食材最好、人手最齐、老板亲自坐镇的日子，评委来了，吃得很满意，给了个好评。

另一家餐馆，某个普通的周二晚上，来了一位没打招呼的密探，吃完走了，没挑出毛病。

哪一条记录更值钱？

答案很明显，但它背后的道理值得说清：两条记录的差别不在评委的水平，不在餐馆的水平，甚至不在那顿饭本身——**差别在于日子是谁挑的。**

《他手无事账》讲的就是这一条：**取样的时刻由谁决定，决定了这份证据还有没有信息。**

为什么今天要紧？

因为“准备”这件事已经可以外包了。

以前，一份精心准备的作品集，本身还能说明一点东西——至少说明这人肯下功夫、有一定水平。现在，准备可以被大规模代做，于是“我准备好了给你看”这类证据整体失效。**凡是你自己选好时刻交出来的东西，都在这个失效名单上：**简历、作品集、项目复盘、事先约定的答辩、你挑好日子请人来看的一切。

剩下唯一还带信息的，是**你没准备的那些时刻**。

这一章有个特别的地方：它发现了一个空格。

书里把这件事拆成三条归属：取样的时刻归谁定、取样的手归谁、记录归谁持有。三条各有两种可能，张成八个格子。

反兴奋剂查了三十年，民航查了几十年，消费信贷也查了几十年——三个行业各自摸索出了自己的那一格，加上毕业生现在的处境，八个格子里占了七个。**剩下一格是空的，四方都没有：他人之手、非本人所选的时刻、无异常也照样入账、并且记在本人可以带走的账本上。**

这一格为什么重要？前三个条件保证这份记录可信；最后一个条件——**归己**——保证它是一个人的资产，而不是一次单方面的监视。少了最后一条，这套东西就变成了对人的管控；有了它，它才是一份可以跟着人走的信用。

读数是什么？

两个：**他择率**（你的证据里，取样时刻不由你决定的占多少）和**归己率**（这些记录里，有多少记在你自己可携带的账本上）。

判据可以直接拿去问任何一份简历：**这个人的能力证据里，有几份的取样时刻不是他自己定的？**

大多数人的答案是零。这不是他们的错——现在几乎没有制度在生产这种记录。而这正是它稀缺的原因：**一份就够，因为几乎没有人有。**

一个小场景。

面试官问：“能不能给我一个例子，是别人在你没准备的时候看到你工作的？”

一个候选人说：不好意思，没有——我的项目都是提前准备的汇报。另一个说：有，去年我们组临时被叫去救一个别人的项目，客户方的技术负责人全程在场，事后他写了一段评价发在群里，我截图留着。

第二个人的那张截图，不是因为内容多好，而是因为**时刻不是他挑的**。

再看两个身边的例子。

突击检查。所有人都知道，通知过的检查和不通知的检查，看到的是两个不同的世界。这件事人人都懂，但很少有人把它反过来用在自己身上——**你的能力证据里，有几份是“不通知”的那一类？**

朋友眼中的你。为什么熟人的评价比自我介绍可信？不是因为熟人更客观，而是因为熟人看到的那些时刻不是你挑的：你排队时的样子、你累的时候怎么说话、你计划落空时的第一反应。**自我介绍是自择时刻的，朋友的印象是他择时刻的**——两者的信息量根本不在一个量级。

三种常见的误读。

第一种：以为这是在鼓励监控。这一章明确反对无对价、全覆盖的观察。它的三个条件里，最要紧的第三个是**归己**：记录必须记在被读者本人可以带走的账本上。没有归己的那一格，它就是监控；有了归己，它才是一份跟着人走的信用。这一格的空缺，正是这一章的核心发现。

第二种：以为“没发现问题”不算成绩。恰恰相反，这一章最反直觉的一条就是：**无异常也要入账**。现行所有检查制度只记录问题、不记录“没问题”，于是一个从没出过事的人和 一个从没被检查过的人，在账面上完全一样。把“无事”变成可累积的资产，是这一格与前七格的关键差别。

第三种：以为可以事后补。补不了。一份事后追认的“当时我表现不错”，它的时刻仍然是你自己挑的。他择记录必须在当时、由对方留下。这就是为什么它稀缺，也是为什么一份就够。

给自己做个小检查。你手上所有能拿出来的能力证据里，有几份的**时刻**不是你自己定的？如果答案是零，这不是你的失败——现在几乎没有制度在生产这种记录。但它意味着，你今天开始留的第一份，就是罕见的。

它和旁边几个概念的区别。

和《脱稿步》的两轴正交，上面已经讲过：一个管从哪儿出发，一个管时刻谁定。

和《峰兑》的关系是包含而非等同：峰时事件天然他择，但他择不必是峰时——一次普通周二的无预告抽查，也是完整的他择记录，而且比事件常见得多。这一点很实用：**等事件可能要等好几年，而他择记录今年就可以开始积累。**

和《零判份》互不相干却常被混为一谈： ρ 说的是产出本身的判别力，他择率说的是这份读数可不可信。一份 ρ 很高的作品集，如果全部是自择时刻交出来的，它的可信度问题一点没解决；反过来，一次他择读数也不会提高你产出的未见率。产出轴与取证轴是两码事。

和《入账之前的技能》在“归己”这一格上咬合：他择记录如果只记在机构账上（像飞行数据那样），那么这一章的第三个条件不满足，八格里的那个空格仍然是空的。

一句话记住它：这一章管的是凭什么信——而唯一的答案是，那个时刻不是你挑的。

怎么长出来：方法、技术与流程

个人层面。

第一，**主动进入不由你择时的场合**：临时顶上的活、别人挑时间的评审、突发状况下的支援、无预告的现场答疑。

第二，**把这类记录留下来并归己**：请对方留一句话、留一份可查证的痕迹，自己保存一份。这件事必须当时做，事后补的都带着自择的味道。

第三，**分清两种自愿**：进入检查库应当是你自愿的、有对价的（换来机会、优先权、可携带的记录）。无对价、被强加的观察不是这一章的处方，是它明确反对的东西。

大学发生学教育：建一个自愿的“检查库”，并把记录还给学生。

这一条是发生学教育里制度性最强的一条，也最需要学校层面的设计，但它的成本比想象的低——它主要是记账方式的改变，不是新增环节。

- **自愿检查库+对价**。学生自愿报名进入一个“可被无预告取样”的名单，报名换来实际的对价：进入更好的项目、优先获得实习推荐、获得一份可携带的记录。**不报名不受任何不利影响**——这一条必须写死，否则它就是强制观察换了个名字。
- **无预告随堂应答**。对库内学生，教师可在任意一次普通课堂上，就他手上正在做的东西做十分钟的当场应答（与《脱稿步》那一章的追问共用同一套动作，但这里的关键变量是**无预告**）。
- **无异常也入账**。这是最容易被忽略、却最要紧的一条：**检查没发现问题，也要记一笔“某年某月某日，无预告取样，无异常”**。现行的所有检查制度都只记录问题，不记录“没问题”——于是一个从没出过事的人，账上一片空白，与从没被检查过的人无法区分。这一格的价值，正在于把“无事”变成可累积的资产。
- **归己的账本**。这些记录必须以学生可携带的形式给他本人（一份带校方签章的、可随简历一起提交的记录），而不是只留在学校系统里。**记录归谁持有，决定了它是资产还是监控**——这是全书归属层的核心，也是这一格与前七格的分界。

教师与教务的动作。

第一，编码规则公开：什么算无预告、什么算无异常、谁来记，全部写在明处，并留一条学生可以申诉与要求撤出的通道。第二，抵抗把它变成抽查纪律的诱惑——它记的是能力证据，不是考勤，两者一旦混同，学生会立刻把它当成敌意的东西并开始规避。

最容易做坏的地方。

假他择：说是无预告，实际上内部提前通气。它会以最快的速度掏空这套制度的全部价值，而且很难查。唯一的防线是编码规则公开+抽样的时间由不参与教学的第三方决定。

原文

他手无事账：生成式人工智能时代大学毕业生核心竞争力的证据落点

——基于反兴奋剂、民航自动化与消费信贷三个参照案例的分析

王德生 德麦国际 Demai International Pte. Ltd.，新加坡

The Other-Hand Clean Ledger: Where the Evidence of Graduate Competitiveness Goes in the Age of Generative AI — An Analysis Based on Three Reference Cases from Anti-Doping, Cockpit Automation and Consumer Credit

Wang Desheng, Demai International Pte. Ltd., Singapore

摘要

目的：回答一道被生成式人工智能改写了前提的问题——大学毕业生的核心竞争力将是什么。既有回答大多给出能力清单，而清单默认能力可以从毕业生交出的产物上读出；生成式 AI 使产物与能力之间的推断关系归零，问题因此变为“当产物不再是证据，能力的证据在哪里”。**方法：**本文选取三个此前二三十年间各自撞上同一问题、且互不引用的领域作为参照案例——反兴奋剂的运动员生物护照、民航驾驶舱自动化后的手动飞行小时、消费信贷中的“信用隐形人”——把三家的回答写成同一句争论，找出三家共同默认而未说出的前提，并用其中一家自己的实证材料（微剂量促红素在护照标记下不被标出）将其推翻；随后给出承重命题、辨别装置与可跨领域通约的读数，并对两组公开文件做预注册的字段审计。**发现：**三家共同默认“存在一处不能被制造的参照物，证据搬到那里就安全”；该前提在三家各自的实证材料里均已被推翻。推翻之后剩下的不是第四处参照物，而是三条三家都用过却无一家当作对象的归属：**取样时刻归谁定、取样之手归谁、记录归谁持有**。三条归属张成八格，三家与毕业生分别占了七格，**有一格四方都空着：**由他人之手、在被读者不选择的时刻取得、结果为无异常、且记入被读者可持有账本的记录。本文把这一格命名为**他手无事账**，并提出：毕业生的核心竞争力不是产物质量、不是自动化之外的残余技能、也不是被记录的厚度，而是他手无事账的存量。两次审计：六所公布“双车道评估”的大学全部以监考与定时定义安全车道，无一以无预告定义，与本文的次序预测一致；美国联邦航空局 2022 年《飞行路径管理》咨询通告的三条预注册预测全部命中，而一条预注册之外的发现——民航早已持有逐次航班的他择读数却以政策不归于个人——正是把本文的承重物从“时刻归属”这一条属性推到“三条归属张成的那一格空位”的材料。第三次审计按预注册去四个职业场域找那一格的反例（会计师事务所同行检查、医师执业再认证、照护机构无预告检查、反兴奋剂机构公开的运动员检测次数）：无一同时满足三个条件，最近的一个差半个条件。

结论：生成式 AI 没有消灭能力，它使自择时刻的证据归零；四个场域各自持有那一格的两个条件而无一持有第三个，这一格空位就是竞争力将要落进的地方。

关键词：他手无事账；他择存量；他择率；自择证据；取样之手；记录归属；运动员生物护照；手动飞行小时；信用隐形人；双车道评估；毕业生就业能力

Abstract

Purpose. This paper addresses a question whose premise has been rewritten by generative AI: what will be the core competitiveness of university graduates? Existing answers list competencies, and lists presuppose that ability can be read off what graduates submit; generative AI reduces the inferential link between product and ability to zero, so the question becomes where the evidence of ability goes once products cease to be evidence. **Method.** Three fields that met the same problem two to three decades earlier and do not cite one another are taken as reference cases: the Athlete Biological Passport in anti-doping, manual-flying hours after cockpit automation, and “credit invisibles” in consumer finance. Their answers are written as one dispute, the premise all three share but never state is identified, and it is refuted with one field’s own empirical material (micro-dose EPO is not flagged by passport markers). A load-bearing proposition, a discriminating grid and a cross-field commensurable ratio are then given, and two sets of public documents are audited under pre-registration. **Findings.** All three fields assume that a reference point exists which cannot be manufactured, so that evidence is safe once relocated there; each field’s own data refute this. What remains is not a fourth reference point but three attributions all three fields used and none treated as objects: *who chooses the sampling moment, whose hand takes the sample, and who holds the record*. The three attributions span eight cells; the three fields and the graduate occupy seven, and one cell is empty for all four: a reading taken by another’s hand, at a moment the person did not choose, showing no anomaly, and entered into a ledger the person holds. We name this cell the **other-hand clean ledger** and propose that graduate competitiveness is not product quality, not residual un-automated skill, not record thickness, but the stock of that ledger. Two audits: all six universities with published two-lane assessment frameworks define the secure lane by supervision and scheduling, none by no-notice sampling, as the ordering prediction requires; three pre-registered predictions on the FAA’s 2022 Flightpath Management Advisory Circular are all met, while one unregistered finding—aviation already holds flight-by-flight other-timed readings but by policy keeps them de-identified—is precisely what pushed the load-bearing object from one attribute (timing) to the empty cell spanned by three attributions. A third pre-registered audit searched four professional domains for a counter-example to the empty cell (CPA firm peer review, physician revalidation, unannounced care-home inspection, publicly posted athlete test counts): none satisfies all three conditions; the

nearest misses by half a condition. **Conclusion.** Generative AI did not abolish ability; it zeroed self-timed evidence. Each of the four domains holds two of the cell's three conditions and none holds the third; that empty cell is where competitiveness will land.

Keywords: other-hand clean ledger; other-timed stock; other-timed ratio; self-timed evidence; Athlete Biological Passport; manual flying hours; credit invisibles; two-lane assessment; graduate employability

1 引言

1.1 问题的改写

“大学毕业生的核心竞争力是什么”在 2023 年之前是一道常规问题，标准答案是一张能力清单：批判性思维、沟通、协作、终身学习，再加一门专业（World Economic Forum, 2020；OECD, 2019）。清单的问题不在内容，而在它默认了一件事——这些能力可以从毕业生交出来的东西上读出。简历、作品集、成绩单、面试作答、试用期报告，每一项都是毕业生在自己选定的时刻、准备完毕之后交付的产物；用人单位读产物，推断能力。

这条推断链在 2023 年之后失效。一份结构完整的报告、一段可运行的代码、一份获奖水准的作品集，任何会使用生成式工具的人都能在一小时内交出。产物与能力之间的推断关系不是变弱，而是趋于零；用人单位不是变得更难判断，而是他们习惯读取的整类证据同时失效（Kofinas et al., 2025）。

于是问题不再是“毕业生应当具备什么能力”，而是“当产物不再是证据，能力的证据在哪里”。这个问题教育学不是第一次遇到，也不是最早遇到。至少三个行业在此前二三十年各自撞上同一堵墙：竞技体育的成绩可由药物制造，反兴奋剂机构被迫重新决定证据放在哪里；民航飞行可由自动驾驶完成，监管者被迫重新决定“飞行员会飞”的证据放在哪里；消费信贷要求在放款之前读出偿付能力，而年轻人手上没有任何可读的产物。三家各自给出了答案，三个答案互不引用，且后来各自出现了同一种失效。

1.2 本文的贡献

本文不把三家当作三个类比，而当作三个参照案例：把它们的答案写成同一句争论，找出三家共同默认而从未说出的前提，用其中一家自己的实证材料将其推翻，然后看剩下什么。贡献有四：

（1）指出三家共享的一条前提——“存在一处不能被制造的参照物”——并证明它在三家各自的实证材料里均已被推翻；（2）提出三条此前被三家各自使用却无一家当作对象的归属——取样时刻归谁定、取样之手归谁、记录归谁持有——并证明三条归属张成的八格里有一格四方皆空，据此给出承重命题“他手无事账”；（3）给出时刻轴的 2×3 辨别格、三归属的八格表、以及可跨领域通约的比率“他择率”与“归己率”，并说明它们与既有主要指标正交；（4）做三次预注册审计：两组公开文件的字段审计，以及对承重物“那一格是空的”这一主张的反例搜索；报告一条改变了承重物的发现，并把最近的反例写进正文。

1.3 结构

第 2 节回顾相邻文献；第 3 节陈述三个参照案例；第 4 节分析三家的争论、共有前提与推翻材料；第 5 节给出命题、辨别装置与读数；第 6 节报告两处来自参照案例的反向约束；

第 7 节给出证伪条件与三次预注册审计；第 8 节与最近的既有说法划界；第 9 节讨论伦理与适用边界；第 10 节是局限与自反；第 11 节给出结论与写死日期的可检验预测。

1.4 术语约定

自择证据：取样时刻可由被读者推迟或提前的能力证据。**他择证据**：取样时刻由读者决定的能力证据；按告知期分为**他择·定时**与**他择·无预告**两档。**他择率**：他择证据份数占全部证据份数的比率。**他择存量**：他择证据中记录为“未见异常”者的累积份数。**自手／他手**：证据由被读者自己提交或记录（自手），还是由第三方执行取样并记录（他手）。**归己／不归己**：记录是否落到被读者名下并可由其持有携带。**他手无事账**：他择·他手·归己、且记录无异常的证据账目；本文的承重物。**归己率**：他手读数中记入被读者可持有账本的比率。**参照物**：在产物之外被选作证据落脚点的位置（基线、残余操作、第三方记录）。**读者／被读者**：读取证据的一方与被读取的一方；本文用这对词而不用“评价者／被评价者”，因为后者预设了评价关系，而他择读数的读者可以是检查官、观察员、数据记录器，不一定在评价。

2 文献回顾

2.1 信号、筛选与代价

学历为什么有价值？信号理论（Spence, 1973）与筛选理论（Arrow, 1973）的回答是：取得学历的成本对能力低者更高，因此学历能分开人。生物学中的残障原理（Zahavi, 1975）给出同一形状：可信的信号必须昂贵到造假者负担不起。这一支文献的前提是信号的**制造成本**与能力负相关；它从不问信号是在谁选定的时刻发出的。Caplan (2018) 把学历价值的大部分归于信号而非人力资本，但同样不区分信号的时刻归属。

2.2 评估安全、双车道与口试

教育领域对生成式 AI 的回应集中在三支。其一，Dawson (2021) 把“学术诚信”（学生该不该）与“评估安全”（能不能作弊）分开，主张后者是设计问题，手段包括身份认证、监考与工具限制。其二，“双车道”评估（Bridgeman & Liu, 2024 ; Corbin, Dawson & Liu, 2025）把评估分为受监考的“安全车道”与允许 AI 的“开放车道”，已被悉尼、奥克兰、巴斯等校写入政策；批评者认为它过于二分（Curtis, 2025）。其三，口试与“同步核验层”（Joughin, 2010 ; Sotiriadou et al., 2020）主张在作品之外由评估者当场提问，近两年的框架给出了三段式口试规程（Frontiers in AI, 2026）。三支的共同点是把注意力放在取样的**条件**（谁控制环境、工具、身份）上，本文将说明取样的**时刻**是另一根轴。

2.3 自动化与残余技能

Bainbridge (1983) 指出自动化把操作者变为监视者，而技能在最需要时已退化。Casner 等 (2014) 在模拟机上发现，长期依赖自动化的航线飞行员手动操纵技能基本保留，退化的是自动化断开时的认知技能。这一支文献是本文第二个参照案例的正主，第 4 节将说明它同时提供了推翻自己参照物的材料。

2.4 检查博弈与不预告观察

博弈论早已把不预告检查形式化（Dresher, 1962 ; Avenhaus, von Stengel & Zamir, 2002）：检查者与被检查者在时刻选择上博弈，最优检查策略是随机化。软件可靠性中的混沌工程（Basiri et al., 2016）在生产系统里不预告注入故障。餐饮卫生等级牌

(Jin & Leslie, 2003) 把不预告检查的结果累积并公开, 使食源性疾病住院下降约五分之一。这一支把不预告当作**检查者的策略或场所的属性**, 无一把它当作被检查者可持有的**证据属性**。社会学中的前台/后台 (Goffman, 1959) 把被观察与否当作空间的属性而非时刻的归属。

2.5 试用期与实习: 用人单位既有的他择

用人单位并非没有他择读数。试用期是一段由用人单位定时、被雇者不能推迟的观察; 实习评价也常包含主管对临时任务的记录。这一支实务没有进入评估文献, 原因是它的结果不可携带: 试用期通过与否只对该单位有效, 实习评价通常是一份自由文本, 没有任何字段记录“哪一次是临时的”。它们是他择读数, 但不累积、不流通——落在表 2 的右列。本文与它们的关系在 8 节与 9.5 节给出。

2.6 缺口

上述六支各自触到“取样时刻由谁定”“取样由谁执行”“记录归谁”三条归属中的一条或两条, 但无一把三条并排、无一问过三条同时成立时会出现什么; 它们也没有一支把“参照物本身可被制造”当作问题。本文补的正是这两处。

3 三个参照案例

3.0 案例选取

三个案例按四条标准选出。其一, **互不引用**: 反兴奋剂文献、人因工程文献与消费信贷文献之间没有交叉引用网, 三家对同一问题的回答因此是独立形成的, 而不是相互抄的。其二, **各自撞过同一堵墙**: 三家都经历了“产物可被外力制造”这一事件, 且都在事件之后重新选择了证据的落点。其三, **各有可清点对象**: 护照样本、航段、账期, 三者都可数, 使第 5 节的读数有取数场。其四, **各有厚的判据体系**: 国际体育仲裁法庭十余年的护照判例、民航事故调查与咨询通告体系、信用评分的监管与诉讼史, 使“对立的答案”不是主张而是有判据的做法。三家分别从三个位置回答——被显露出来的形状、一段不可替代的路径、证据得以成立的场——这一点在表 1 最后一栏标出; 同一位置上的两家只能靠判谁对来化解冲突, 那是辩论而不是本文要的并排。

3.1 反兴奋剂: 证据从成绩搬进基线

竞技体育最早的证据是成绩本身。药物出现后, 证据被搬进尿样。到 2000 年代中期尿检失效: 促红素的检测窗口只有几十小时, 运动员在窗口之外用药, 比赛时的样本干净。

2009 年, 世界反兴奋剂机构把证据再搬一次, 搬进一个此前不存在的位置: **运动员自己的历史**。运动员生物护照不再问“这份样本里有没有药”, 而问“这份样本与此人过去十几份样本的分布是否一致”。网织红细胞百分比、血红蛋白浓度被逐次记录成纵向轨迹, 判定依据是**偏离自身基线的程度**(WADA, 2009 及后续版本)。国际体育仲裁法庭自 2010 年起确认仅凭纵向轮廓、未检出任何药物的处罚成立。护照的判定不是自动的: 软件先按贝叶斯模型给出该样本落在此人预期范围之外的概率, 超过阈值的“非典型护照结果”交由三名独立专家匿名评审, 专家须排除海拔训练、疾病、脱水等生理解释, 三人一致认为“极可能是禁用方法所致”才立案。也就是说, 护照把“证据”重新定义为**一个人与他自己的偏差**, 而不是一个人与某个外部标准的偏差; 一名血红蛋白偏高但一向如此的运动员不会被标出, 一名数值在群体正常范围内但相对自己突然变化的运动员会被标出。这个设计的前提是自己的历史无法被自己改写。

护照还带来一个此前不存在的读数方向：**进步本身成为可疑对象**。在产物型证据时代，成绩提高是能力的证明；在护照时代，一段没有生理来路的提高是立案的触发。这一点与本文第 5 节的判据有直接关系——当产物可造，“越来越好”不再是无条件的正面信号，“在谁选的时刻被读到”才是。

这一步的含义是：证据从产物（成绩、样本中的药）搬到了参照物（此人自己的历史），并默认该参照物不能被制造——生理基线是自己的，药物只能使人偏离它。与此同时出现的“行踪规则”要求注册检查库中的运动员每天报出一个 60 分钟时间窗，检查官可在窗内不预告地取样；一年三次错过等同一次阳性。行踪规则不是护照的附件，而是护照运转的条件：纵向轨迹要成立，取样时刻就不能由运动员决定。

3.2 民航：证据从飞行搬进手动小时

1980 年代以后驾驶舱自动化逐年上升，到 2000 年代，一段典型航线飞行中飞行员手动操纵的时间只有起降前后的几分钟。问题出在自动化失效的时刻：法航 447（2009）皮托管结冰、自动驾驶断开，机组四分钟内未能改出可恢复的失速（BEA，2012）；韩亚 214（2013）自动油门处于机组未察觉的模式，进近速度持续衰减（NTSB，2014）。两份调查指向同一点：飞行员会不会飞，不再能从他飞过的航段读出——那些航段是自动驾驶飞的。

监管者把证据搬到新位置：**未被自动化替代的那段操作**。美国联邦航空局 2013 年发布安全警示 SAFO 13002《手动飞行操作》，要求航空公司鼓励飞行员在适当条件下手动飞行，理由是持续使用自动驾驶可能导致手动技能退化；2017 年 SAFO 17007 再次重申。手动飞行小时、模拟机中不带自动化的科目、航线检查中的手动进近，成为“此人还会飞”的证据。搬动的形状与第一家相同，且同样默认参照物不能被制造——手动飞过的小时就是手动飞过的。

值得记下的是监管者在 2013 年前后拿到的那份行业报告。联邦航空局的绩效评审委员会与商业航空安全团队 2013 年联合发布《飞行路径管理系统的运行使用》，汇总了二十六起事故与重大事件、九千余份自愿报告与航线观察数据，结论之一是：机组对自动化的依赖不只是技能问题，而是**知道飞机现在在做什么**的问题——飞行员在自动化模式切换时常常不知道飞机处于哪个模式。这份报告已经把问题从“手会不会动”移到了“脑子知不知道”，但监管文件的处方仍然落在“多手动飞”上。这是第 4 节要用到的一处错位：参照物记的是手，缺的是脑。

3.3 消费信贷：证据从收入搬进记录的厚度

一个人能否按时还款，放款人须在放款前判断。收入证明是产物型证据——可伪造，也可在借款后消失。1950 年代起，美国消费信贷业逐步把证据搬进第三方记录：信用档案。每一笔账户的开立、每一期还款的准时与逾期，由三家征信机构逐月记录，评分模型从记录中读出一个数。搬动逻辑仍然相同：产物可造，记录不能——它是别人替你记的。

这一家暴露出另外两家未正面遇到的问题：**没有记录的人**。美国消费者金融保护局 2015 年的统计：约 2600 万成年人在三大征信机构没有任何档案，另有约 1900 万因记录太少或太旧而无法评分，两类合计约占成年人口五分之一，且高度集中在 25 岁以下（CFPB，2015）。这些人不是信用差，而是**信用不可见**。一名应届毕业生在这套系统中的位置正是“薄档案”。

这一家还出现了另外两家未出现的事：既然证据是记录厚度，市场上出现了**制造厚度的产品**——信用建设贷款（按月向锁定账户存钱，每期被报送为准时还款）、授权用户（挂到

他人老账户上，其历史立刻出现在自己档案里）、以及曾公开交易的”交易线出租”。评分公司在 2008 年模型中专门加入过滤以对付最后一种（FICO, 2008）。

信贷这一家还有一处另外两家没有的制度特征：**正向记录**。反兴奋剂只记异常样本，民航只记事故与偏差；信贷从 1970 年代起同时记录准时还款——一次”没出事”也进档案，并且越多越好。正因为如此，信贷是三家里唯一把”无事”累积成资产的一家，也是三家里唯一出现了制造”无事”的市场的一家。这两件事不是巧合：**一旦无事被累积成资产，无事就可以被制造**。这一点在第 5.5 节与第 9.3 节都会回来。

4 分析：争论、共有前提与推翻

4.1 三家争的是同一件事

三家的争论可写成同一句：**当产物可以被外力制造，能力的证据该放在哪里**。三个答案见表 1。

表 1 三个参照案例的证据落点

参照案例	证据搬到了哪里	该位置被认为不能被制造的理由	回答的维度
反兴奋剂	此人自己的纵向基线	生理基线是自己的，药只能使你偏离它	被显露出来的形状
民航	自动化之外的残余操作	手动飞过的小时就是手动飞过的	一段不可替代的路径
消费信贷	第三方逐月记的档案	别人替你记的，你改不了	证据得以成立的场

三家从三个位置回答；它们的冲突不能靠”判谁对”化解，因为各自的问题结构使另外两家的答案对自己不可见：反兴奋剂不关心你有没有手动飞过，民航不关心你有没有基线，信贷不关心你偏没偏离自己。三家的答案在各自领域都是标准做法。把它们并排放好之后要问的不是哪一个更好，而是它们底下共同站着什么。

4.2 共有前提

三家的答案有一个共同形状：证据从产物搬到某个参照物，**并默认该参照物不能被制造**。这不是三家各自的立场——三家都没把它当立场说出，它是三家动手之前已经站在上面的地面。念给三家听，三家都会说：这还用说吗。

写实到本题：

存在一处不能被制造的参照物；把能力证据搬到那里，产物的可制造性就被绕开了。

它有一个更靠上的结构层形式：**能力可以从某一个位置单独读出——只要位置选对，一处读数就够**。三家争的是哪一处，没有争的是”有这么一处”。

这条前提放到毕业生身上，正是 2023 年以来全部对策的地面：把考试搬回考场（参照物：监考下的作答），把评价搬进过程（参照物：草稿、版本、思考记录），把证据搬进实习与推荐信（参照物：第三方记录）。三种对策各对应三家的一种，各自默认自己搬去的地方不能被制造。

4.3 推翻材料：来自第一家自己

推翻该前提的材料不必外求，第一家自己早已持有，只是它在自己的领域里回答的是另一个问题。Ashenden 等（2011）给十名受试者按微剂量方案注射促红素——每周两次、剂

量只有常规方案的几分之一，持续十二周——然后用护照的判定模型逐份读取血样。结果：**护照的现行标记未标出任何一名受试者**，而其血红蛋白总量确实上升。该文在自己的领域里回答的是“护照灵敏度够不够”，结论是不够、需要更多标记与更密取样。没有人把它当作关于那条前提的证据。

而它说的正是：**参照物也是可以被制造的**。微剂量方案做的不是使运动员偏离基线，而是使基线本身在判定阈值之内缓慢上移——一条被制造出来的基线。这件材料是一个观察结果而非主张，是反兴奋剂自己的实验数据；它推翻的是反兴奋剂自己站的那块地。

4.4 三处同形的裂缝

第一家的参照物被证明可造，另外两家的参照物立刻露出同一条裂缝，证据同样来自它们自己：

- **民航**：Casner 等（2014）发现长期依赖自动化的飞行员手动操纵技能基本保留，退化的是自动化断开时判断飞机在哪里、处于什么模式、下一步做什么的认知技能。“手动飞行小时”这个参照物记的是保留得最好的那部分，漏掉的恰是事故中缺的那部分；同时手动小时可在模拟机中积累，而模拟机是自动化程度更高的设备。
- **信贷**：信用建设贷款、授权用户、交易线出租构成一个公开的、以制造档案厚度为目的的产业；评分公司加过滤，说明它自己承认参照物在被制造。

三处裂缝的形状相同：**每一处被选作证据落脚点的参照物，一旦成为落脚点，就成为下一轮制造的目标**。这不是三家做得不够好，而是“搬到一处不能被制造的地方”本身不成立。到这里，三家都不能退回自己的阵地：反兴奋剂不能说“把基线记得更密”——它自己的数据说更密也标不出微剂量；民航不能说“多记手动小时”——它自己的研究说手动小时记的不是缺的那部分；信贷不能说“把档案记得更厚”——它自己的市场在卖厚度。

4.5 剩下的不是第四处地方，而是一条维度

前提推翻之后，最省事的反应是找第四处地方；这条路走不通，理由已在 4.4。要看的是三家在参照物失效之后各自**实际**又做了什么。

反兴奋剂在护照失灵后没有再搬地方，它做的是把**取样时刻**从运动员手中拿走——行踪规则、无预告检查、2010 年代后期从定期取样转向情报驱动的靶向取样（Lauritzen, 2023）。判定依据仍是那条轨迹，变的是轨迹上每一个点由谁决定何时取。民航在手动小时之外真正保留的证据是**航线检查**：检查员不预告地坐进驾驶舱，在这一段真实航线上看此人怎么飞。信贷评分中真正抗制造的部分，行业称为“经历过周期的档案”：一个账户在 2008 年下行期的表现比其扩张期十年准时还款更值钱，因为**下行期不是借款人选的**。

三家各自走到的地方从外面看是三件不相干的事——行踪规则、航线检查、周期表现——但它们是同一件事：**证据的价值不由它落在哪个位置决定，而由取样时刻是否由读者、而非被读者选定决定**。一份运动员自己选好时间提交的血样、一段飞行员自己选好科目练出的手动飞行、一条借款人在自己选好的宽松年份建立的档案，三者都是**自择证据**；一份 60 分钟窗口内被敲门取走的血样、一段检查员不预告坐进来看的进近、一段谁也没选的下行期里的还款，三者是**他择证据**。

这条维度此前不是没被用过，而是从未被当作一样东西。三家各自用了它，各自用自己领域的名字称呼它，各自把它当程序细节。它落在三家账本的缝里：护照记数值不记时刻由谁定，航线检查记通过不通过不记是否预告，档案记准时不记那一年好不好过。

必须排除一种更省事的读法：三家走到的地方不就是”第三方”吗？不是。信贷的档案从一开始就是第三方记的，它照样被制造——因为借款人虽然不能改记录，却能决定**什么时候**开一个账户、什么时候还一笔款；第三方记的是时刻由本人定的事件。模拟机训练也是第三方执行的，飞行员照样能选科目、选时段。第三方决定了**谁记**，没有决定**何时取**；而三家在参照物失效后各自抓住的，恰恰是何时取。把”他择”压缩成”第三方”，会把信贷档案与行踪检查放进同一格，而它们在三家的实际历史里正是被分开的两样东西。

还须排除另一种读法：他择不就是”随机”吗？也不是。行踪规则下的取样常常不是随机的，而是情报驱动的靶向取样——检查官专挑赛前几周、专挑训练营；航线检查也常常安排在困难机场。他择的定义只要求时刻不由被读者决定，不要求读者的选择无规律。随机是他择的一种实现，不是它的定义。

4.6 三家的次序为什么相同

三家从产物到参照物、再到时刻归属，走的是同一条次序，且各自用了大约一代人的时间（表 1a）。

表 1a 三个参照案例的证据迁移次序

阶段	反兴奋剂	民航	消费信贷
产物即证据	成绩（—1960s）	完成的飞行（—1980s）	收入证明（—1950s）
产物中的痕迹	尿检阳性（1968—）	事故调查（持续）	抵押物（持续）
参照物	纵向基线／护照（2009）	手动飞行小时（2013）	第三方档案（1970s—）
参照物被制造	微量方案（2011 实证）	模拟机小时；手记脑不记（2014 实证）	信用建设产品、交易线出租（2000s）
时刻归属	行踪规则（2009，与护照同时）、靶向取样（2010s 后期）	航线检查（一直存在，被重新倚重）	周期表现（2008 后）
存量是否归人	归人，只记异常	持有而不归人（7.5）	归人，正向记录

次序相同不是巧合，而是同一条逻辑的三次执行：证据落到哪里，制造就跟到哪里；能被制造的位置终究都是被读者可以选时刻的位置；剩下的只有读者选时刻的那些点。三家花了一代人各自走到这一步，并且都没有把这一步命名——这正是教育侧可以省下的那一代人。

4.7 第二条与第三条归属：取样之手与记录归谁

4.5 抓住的是时刻归谁定。但三家在参照物失效之后的实际做法里，还藏着另外两条归属，同样各自用了、各自没当成对象。

第一条是**取样之手归谁**。护照的样本由检查官取，不由运动员交；航线检查由检查员坐进来看，飞行数据由记录器记，不由飞行员填；信贷档案由放款机构报送，不由借款人自述。三家在这一点上一致：证据落脚点失效之后，它们都把取样从被读者手里拿走了。信贷早就这样做，但它的问题在时刻（4.5）；反兴奋剂与民航后来也这样做，把时刻也拿走了。

第二条是**记录归谁持有**。这一条三家分开了。信贷把记录归到借款人名下，借款人可以调阅、争议、携带，且正向记录——“无事”入账。反兴奋剂把护照归到运动员名下，但只记非典型结果，一次清白的敲门不留痕。民航的飞行数据是每一次航班都记的、不由飞行员

选时刻的、不由飞行员之手的——它同时满足前两条归属——**但以明文政策去除身份、不得用于对个人的判定**（7.5）。

把三条归属并排，得到八格（表 1b）。三家与毕业生各自占了哪些格，是可以逐一查出的事实，不是构造。

表 1b 三条归属张成的八格及其占据者

时刻	手	账	占据者
自择	自手	归己	作品集、简历、自选项目——毕业生的目标格
自择	自手	不归己	一次性自评问卷、无记录的练习
自择	他手	归己	信贷档案 （准时还款入账，事件时刻由借款人定）；可预约的考试与面试
自择	他手	不归己	不进档案的实习汇报、内部测评
他择	自手	归己	应临时要求自报的材料（限时自述）
他择	自手	不归己	临时口头问答无记录
他择	他手	不归己	民航飞行数据（去身份） ； 卫生等级牌 （归场所不归人）；闭卷考试成绩单归己，但只记分数不记”无事”
他择	他手	归己且记无事	空 ——反兴奋剂在此格只记异常；民航在此格不归人；信贷在此格时刻自择；毕业生账本无此字段

最后一行是本文的承重物。它不是一个读数、不是一条属性，而是一类**四方账本上都不存在的记录**：由他人之手、在被读者不选择的时刻取得、结果为无异常、且记入被读者可持有的账本。四个场域各自持有它的两个条件，无一持有第三个——反兴奋剂有他择与他手，无”无事归己”；民航有他择与他手，无归己；信贷有他手与无事归己，无他择；毕业生三者皆无。**删掉这一格，四家的账各自都记得平**——这正是它此前从未被当作一样东西的原因。本文把它命名为**他手无事账**。

4.5 的”他择”仍然成立，它是这一格的三个条件之一，也是三家最后才各自抓住的那一个；但以它为承重物会把信贷档案（自择·他手·归己）与行踪检查（他择·他手·只记异常）当作同一件事的两个程度，而它们在表 1b 里隔着两格。承重物必须落在那一格空位上，不能落在通向它的某一条轴上。

5 命题、辨别装置与读数

5.1 承重命题

生成式 AI 时代，大学毕业生的核心竞争力不是他交出来的产物的质量（第一家的划界：产物可造），不是自动化之外的残余技能（第二家：残余可造，且残余不是缺的那部分），也不是被第三方记录的厚度（第三家：厚度可造），而是他手

无事账的存量——由他人之手、在本人不选择的时刻取得、结果为无异常、且记入本人可持有账本的记录，累积起来的那个量。

三重否定各对应一家；三家的证据合起来只能得出它，而三家单独都到不了它。原因不在三家的疏忽，而在它们各自的核心问题使表 1b 最后一行的第三个条件对自己成为不必要：反兴奋剂的核心问题是抓人，抓不到的读数不记，故不会给一次清白的敲门入账；民航的核心问题是安全，而安全文化要求飞行数据不得归人，否则无人再报告；信贷的核心问题是定价违约，它记的是借款人自己发起的账户，不会把”这一期是不是他选的”记成字段。三家各缺一个条件，缺的是三个不同的条件。

5.2 辨别装置：时刻轴的 2×3 格

表 1b 给的是三条归属的八格；本节把其中最难判的一条——时刻归属——展开成三档，与”是否累积”交叉，用于逐项编码。第一根轴分三档：取样时刻由本人定（自择）；由他人定但时刻公布（他择·定时）；由他人定且不预告（他择·无预告）。第二根轴：读数是否被累积成存量。

表 2 能力证据的 2×3 辨别格

	读数被累积成存量	读数不被累积（只在异常时记一笔）
自择	目标格：作品集、成绩单、简历、准备好的面试作答、自选题的比赛。全部形式审查都能通过；AI 使它短期看起来更好，且制度已开始把它移出证据栏（7.4）；一份 AI 造的作品集第一次失效时不产生信号；评价越正规化，这一格越厚	自选时段的一次性认证考试。过了是一张证，没过不留痕
他择·定时	学期闭卷考试、定期考核、标准化工作样本测验。确实累积，但候选人可为公布的时刻准备，AI 可代练可准备的部分。双车道”安全车道”与民航”定期检查”现在所在的格	比赛时的反兴奋剂检查：预告的、只记阳性
他择·无预告	本文答案所在：无预告的现场任务、检查员坐进来的那一段、不由本人选的困难期表现——当它们被累积并记为”无异常”时。这一格在毕业生账本上目前不存在；民航持有它（每次航班的飞行数据）但以政策不归于个人（7.5）	三家实际所在的格：行踪检查、航线检查、周期表现——存在，但只在出事时记一笔

目标格的三条签名齐全：短期指标改善（AI 使自择产物更漂亮）；失效不产生信号（自择产物失效时不知道是它）；自我加固（评价越正规，对自择产物要求越高，自择产物越厚）。这一格不是一眼能看出的坏，它是所有人都在投入的那一格——正因为它通过全部形式审查。

5.3 零情态词判据

在这个人的能力证据里，有几份的取样时刻不是他自己定的？

再配一句问表 1b 最后一行的：这些取样时刻不由他定的读数里，有几份是别人的手取的、记的是无事、又记在他可以带走的账本上？两句合起来就是他手无事账的清点法。第一句可以拿去问一份简历、一套招聘流程、一本成绩单、一份实习评价表，不涉及任何人的判断。它之所以不含「应当」「真正」「充分」这类词，是因为含了这些词的判据只能问人，而问人得到的是一个判断；本文要的是一句可以拿去问流程文件与表格字段的话——答案是一个数，且两个人去问会得到同一个数。它在五个场景中的答案互不相同：（1）典型应届简历：零，每项都是自己选了时间做完再写上；（2）带有”航线检查”式实习评价的简

历：一到两份——若实习单位有不预告抽查制度且写进评价；（3）注册检查库中的运动员：一年十几到几十份，且每份时刻都不由本人定——此数不是从命题推出，而是从行踪规则的实际执行频率查出；（4）经历过 2020 或 2022 年下行期的信贷档案：那几个月的每一期是他择的，行业评分实践确实给予更高权重——亦是查出的；（5）“双车道”大学的学生：安全车道考试由校方定时，但时刻公布——落在他择·定时格。

5.4 读数：他择率

他择率 = 取样时刻不由本人决定的读数份数 ÷ 该人全部能力读数份数。

三个领域量纲不同、比率相同：反兴奋剂里分子是无预告取样份数，分母是全部样本；民航里分子是航线检查与无预告评估的航段，分母是全部记录为该人操纵的航段；信贷里分子是落在非本人选择的困难期内的账期，分母是全部账期。毕业生身上，分子是不预告的现场任务、他人出题的即时作答、非自选时段的真实工作片段；分母是全部被拿来评价的证据项。分子按 5.2 第一根轴的后两档分栏记录，不合并。

四条性质：**无量纲；可事后清点**（每份证据的取样时刻由谁定是一个可从流程文件核对的事实）；**把一个印象变成一个数；与既有主要指标正交**。最后一条最要紧：一名成绩全优、作品集获奖、面试作答无懈可击的毕业生，他择率可以为零；一名成绩中等、但在实习单位三次不预告的现场事故处理中都在场且无异常的毕业生，他择率不为零。两数互不推出。

归己率 = 他手读数中记入被读者可持有账本的份数 ÷ 全部他手读数份数。它是表 1b 第三条归属的读数：民航的归己率按政策为零，反兴奋剂的归己率对异常为一、对无事为零，信贷的归己率接近一。他择率与归己率相乘，得到落进最后一行那一格的份额。

5.5 存量

他择率与归己率是比率，他手无事账是那一格里记录本身的累积。四个场域现在都不累积它：护照不记一次清白的敲门，飞行数据不落到飞行员名下，档案记无事却只记自择时刻的无事，毕业生的表格里连“取样时刻由谁定”的字段都没有。删掉这一格，它不是变得难测，而是不存在——四家的账各自照样记平。这是本文认为自己在指出一类新记录而非给旧栏目一种新算法的依据。离它最近的两个反例已在表 1b 里标出：卫生等级牌与民航飞行数据都落在它上面一行（他择·他手·不归己），差的是同一个条件。

5.6 编码规则与一个算例

要使他择率可事后清点，须有一套判定“这一份证据的取样时刻由谁定”的编码规则。本文给出三条：

（1）**时刻可由被读者推迟或提前** → 自择。提交截止日之前的任何作业、可预约时段的考试、可选日期的面试，均属此列；截止日是上限不是时刻。（2）**时刻由读者定且事先告知，告知期长于被读者可用工具准备所需的最短时间** → 他择·定时。一次提前两周通知的闭卷考试属此列；告知期与准备时间的比较须写在机构规则里，本文不给通用阈值。（3）**时刻由读者定且告知期短于准备所需最短时间，或不告知** → 他择·无预告。行踪窗内的取样、当场追问、临时叫上去处理一个真问题属此列。

一个算例。某应届生的评价材料共 12 项：成绩单上 8 门课（其中 5 门以课程作业结课、3 门闭卷考试，考试日期开学即公布）；一份作品集；一次实习，实习评价由主管填写，评价中记有两次“临时被叫去处理客户投诉”的表现；一次结构化面试（面试题固定、

时间可约)。按上述规则:5 门作业、作品集、面试为自择(7 项);3 门闭卷考试为他择·定时(3 项);实习评价中的两次临时处理为他择·无预告(2 项)。他择率=5/12≈0.42,其中无预告档 2/12≈0.17。若同一人成绩单全优而实习评价中没有那两次临时处理,他择率降到 3/12=0.25、无预告档为零——成绩单没有变,这个数变了。这就是 5.4 所说的正交。

算例还暴露一处编码难点:实习评价那两次”临时处理”之所以能被计入,是因为主管碰巧写了下来;多数实习评价表没有这个字段。他手无事账之所以在毕业生账本上不存在,直接原因就是表格里没有这一栏——不是没发生,是没处记;而那两次临时处理即使被写下,也只是他择·他手,主管写的是评语不是”无事”,评价表归单位不归学生——三个条件仍只占两个。

5.7 三家与毕业生的读数对照

表 3 他择率在四个场域的取数

场域	分子(他择读数)	其中无预告档	分母(全部读数)	取样之手	记录是否归己且记无事	缺的条件
反兴奋剂	赛内检查+赛外检查样本	行踪窗内取样、靶向取样	该运动员全部样本	他手	归己,只记异常	无事不入账
民航	定期检查、航线检查航段	不预告的航线检查;每次航班的飞行数据	记录为该人操纵的全部航段	他手	不归己(去身份)	归己
消费信贷	落在非本人选择的困难期内的账期	下行期的每一期	全部账期	他手	归己,记无事	他择
毕业生	闭卷考试、临时现场任务、他人出题的即时作答	临时现场任务、当场追问	全部被用于评价的证据项	多为自手	无此字段	三者皆无

表 3 最后一栏是本文全部论证的落点:四个场域各缺一个条件,且缺的互不相同。这也解释了 3.3 末尾那件事:只有信贷把”无事”记成资产,也只有信贷出现了制造”无事”的市场——因为信贷缺的是他择,而自择的无事最容易造。他手无事账一旦设立,它的反噬(9.3)不会是造无事,而是造”他择”——把预告的时刻申报成无预告。

5.8 命题说的是证据,不是能力

须防一种误读:本文没有说”他手无事账就是能力”,也没有说”能力已经不重要”。命题的对象是能力的证据落在哪里。能力仍然存在于人身上,仍然由学习、训练与经验形成;变的是能力被读出的位置。一个用生成式工具写出报告的毕业生可能确实理解报告的内容,也可能完全不理解——产物对这两种情形不再有分辨力,这是全部问题所在。他择读数之所以有分辨力,不是因为它更接近”真实能力”这样的形而上对象,而是因为它排除了被读者为该时刻准备的可能。当准备本身可以外包,未准备的时刻是唯一还带信息的时刻。

由此也可以看出本文与能力清单派的关系:清单派说的每一项能力都可能是对的,本文只是指出,无论清单上写的是是什么,它的证据若来自自择时刻,在 2023 年之后都读不出来了。清单没有错,清单的证据栏空了。

6 来自参照案例的两处反向约束

两家反过来削掉了本文原本要主张的东西;两处都动了命题的适用范围或价值排序。

6.1 来自反兴奋剂：行踪规则的代价。 本文原本要写：“他手无事账应当成为所有毕业生评价的主读数。”第一家不允许。行踪规则运行十余年，代价清楚：注册检查库中的运动员每天交出一小时行踪，该规则在欧洲人权法院被挑战，在运动员工会被反复反对；它之所以被容忍，是因为进入检查库是自愿的、有对价的（参赛资格）、且范围受限（精英层）。把这套东西不加限制地推到毕业生身上，得到的是无对价、全覆盖的被不预告观察的义务。**削后的命题：**他手无事账只对**自愿进入某个“检查库”**的人成立——报名参加带无预告成分的项目、接受带抽查条款的实习、加入自己知道会被不预告叫上去的岗位。范围从“所有毕业生”缩到“自愿进入检查库的毕业生”。

6.2 来自信贷：薄档案的结构。 本文原本要写：“他手无事账衡量一个人真实的能力。”第三家不允许。信用不可见的五分之一人口不是能力差，而是**没有机会被记**。他择读数同样只在有人愿意不预告地读你的地方产生；大公司实习生一年可攒下多次，无抽查制度的小单位或自己家里的学生一次也攒不下。**削后的命题：**他手无事账衡量的不是能力，而是**曝露**——一个人被允许在别人选的时刻被读了多少次。它与能力相关，但首先记录的是机会的分布。本文不能把它读成“谁更强”，只能读成“谁被在他择时刻读过”。

两条的后果具体：第一条使本文的处方只能是“建检查库”而非“改评价标准”；第二条使读数在用于任何分配前都必须先问“此人有没有被读的机会”。

6.3 两条约束的合读

两条约束合起来，把他手无事账从“一个人的属性”改成了“一段关系的属性”：它需要一个自愿进入的位置（6.1），也需要一个愿意在那个位置不预告地读你的读者（6.2）。没有位置，读数不合法；没有读者，读数不产生。这使它在分配上与学历不同：学历可以由个人独自积累，他手无事账只能在关系中积累——它的三个条件里有两个（他择、他手）在读者手里，只有第三个（归己）是给被读者的。若命题成立，毕业生之间新的不平等将不在学校排名上，而在谁能进入有读者的位置——这条推论本文不喜欢，但它是从两条约束里直接掉出来的，不能删。

7 证伪条件与两次预注册审计

7.1 最强竞争解释

其实不就是把考试搬回考场吗。 监考下的作答也是他择时刻，双车道的安全车道正是此。若该解释成立，本文只是给“回到闭卷”起了新名字。分离点在 5.2 的第一根轴：考场是他择·定时，本文的核心是他择·无预告。竞争解释预测安全车道足够、效度稳定；本文预测安全车道的效度将被可穿戴设备、微型耳机、代考侵蚀，教育侧会被逼向第三档，**次序与反兴奋剂相同**——先定时他择，失效后才无预告他择。

7.2 六条证伪条件

- 1. 机制条件：**控制学业成绩与作品集评分后，若“他择·无预告读数份数”与入职后十二个月绩效之间不存在剂量-反应关系，则本命题为伪——届时竞争力应归因于自择产物质量。数据：任一同时保存招聘评分与入职绩效的机构均有。
- 2. 正向次序条件：**采用双车道的大学中，“安全车道加入无预告成分”的时间应晚于“安全车道定义为定时监考”的时间、早于“安全车道被宣布不足”的时间。若任一校一开始即以无预告取样定义安全车道，次序命题为伪。

3. **工作样本效度条件**：人事选拔文献的下一轮元分析中，若标准化工作样本测验的效度在 2023 年后样本中不低于 2023 年前（对照 Schmidt & Hunter, 1998），则本文关于自择成分失效的预测为伪。
4. **信号方向条件**：若 2023 年后招聘方转向更长、更难的自择产物而非现场任务，则信号理论成立、本文为伪。
5. **正交条件**：任一同时有作品集评分与他择读数的毕业生样本中，若两者相关高于 0.6，则他择率不是独立读数。
6. **存量条件**：若在 11.3 写死的日期前没有任何机构把他择读数累积成被评价者可持有的记录，本文关于“这一栏将出现”的预测为伪。
7. **空格条件**（改判后新增）：表 1b 最后一行的“空”是一个可被单件反例推翻的主张。若任一有公开记录的场域找到一类记录同时满足他择·他手·归己且记无事，且该记录可由被读者携带流通，则那一格不空，本文的承重物退回为“他择存量”这条属性。本文在 7.8 按预注册对四个场域做了第一次搜索。

7.3 审计方法

三次审计均遵循同一纪律：假设、样本规则、判负条款、不利结果处置与停止规则在阅读材料之前写死（附录 A 为预注册原文）；材料为公开文件，可复算；结果无论有利不利照抄进正文；代理变量明写。

7.4 审计一：教育侧现在停在某一格

- **样本**：一次检索（检索词：two-lane approach assessment secured lane open lane universities）返回的前六个大学官方或准官方页面：悉尼大学（经第三方汇总页转述）、澳大利亚天主教大学、奥克兰大学、巴斯大学（两页）、贝尔法斯特女王大学。不补样、不换词。
- **字段**：安全车道的定义依据——(a) 监考/到场；(b) 校方定时；(c) 无预告取样。
- **判负条款**：≥2 所以 (c) 定义 → 次序命题为伪，7.2 第 2 条撤回。
- **结果**（附录 B）：六所全部以 (a)+(b) 定义，无一以 (c) 定义。悉尼：安全车道=到场、受监考的学习评估，仅允许“能通过监考可靠控制 AI 使用”的评估类型；天主教大学：安全/开放两车道；奥克兰：AI 仅在“受控评估（第一车道）”“中可被限制，2025 年 11 月修订程序；巴斯：2026/27 学年起以封闭/开放两车道取代原三色分级；贝尔法斯特：安全评估=到场监考。**判负条款未触发。**

该审计逼出两件本文原本没有的东西：其一，贝尔法斯特那一页自述可穿戴设备与微型耳机使监考“越来越困难”，其作者估计现有评估的多数可被中等技术水平的新手用 AI 满意完成——本文预测的侵蚀已被安全车道的推行者自己写下，教育侧已站在反兴奋剂 2000 年代中期的位置。其二，悉尼把开放车道默认值从“禁止 AI”改为“允许 AI，除非教师退出”，并要求每个学位在 2025 年底前提交评估计划——自择产物那一格被制度性地放弃了作为学习证据的地位，比本文的目标格判断走得更远；本文据此把目标格第一条签名改写为“AI 使自择产物短期更好，且制度已开始把它移出证据栏”。

代理变量与限制：以“安全车道定义”代理“教育侧证据落点”；六所均为英语世界大学，不含中国高校；一次检索前六位有排序偏差。该审计证明的只是次序命题在这一小片样本上未被推翻。

7.4.1 三种对策的对应读法

审计一还允许把 2023 年以来教育侧的三种对策逐一放进表 2。**回考场**（安全车道）落在他择·定时·累积格：时刻由校方定但公布，候选人可为之准备，其可准备部分可由生成工具代练；贝尔法斯特那一页自述的侵蚀说明这一格正在重演反兴奋剂赛内检查的失效。**进过程**（草稿、版本、反思日志）落在自择·累积格：学生决定何时保存一个版本，过程本身可被生成——该支文献自己承认过程材料不能免于伪造。**进实习与推荐信**（第三方记录）落在自择·累积格的另一侧：实习的时段、项目、汇报由学生与单位协商而定，推荐信在学生请求的时刻写出；它是第三方记的，但记的是时刻由本人定的事件，与 4.5 排除的”第三方”读法同形。三种对策分别对应三家的三种搬动，且分别落在三家已被证明可造的那三处；没有一种落在他择·无预告格。这不是三种对策的设计者疏忽，而是他们与三家站在同一块地面上。

7.5 审计二：民航侧的证据落点有没有移动

- **材料**：美国联邦航空局 2022 年 11 月 21 日发布的咨询通告 AC 120-123《飞行路径管理》全文（46 页），SAFO 13002 与 SAFO 17007 的正式后继文件。只读这一份。
- **三条预测（写在打开文件之前）**：(P1) 文件把手动飞行从”鼓励在航线上练”（自择）推进到”由教员/检查员定期评估”（他择·定时）——预期出现；(P2) 文件把不由飞行员选择时刻的数据源（飞行数据监控 FOQA、航线观察 LOSA）列为证据来源——预期出现；(P3) 出现”无预告”取样——按次序命题预期**不出现**。
- **判负条款**：P1 不出现 → 第二家仍停在残余路径格，本文”三家都在向他择移动”在民航侧为伪；P3 出现 → 民航已跳过定时档，次序命题为伪。
- **结果**（附录 C）：P1 出现——文件写明飞行员”应被定期评估其手动飞行熟练度”，并为教员/检查员单列考虑事项；P2 出现——FOQA、LOSA、ASAP 被列为训练与政策的反馈来源；P3 不出现——全文零次”无预告”。三条全部命中，判负条款未触发。

一条预注册之外的发现，照记，并说明它改了什么：文件把 FOQA 与 LOSA 只当作**训练大纲与运营政策的反馈数据**，从未当作**个人的能力证据**——这不是疏漏，而是民航行业几十年的明文安排：飞行数据监控按协议去除个人身份、由工会与公司共管、不得用于对个人的惩戒。也就是说，民航早已持有不由飞行员选时刻、不由飞行员之手的逐次航班记录，并且**有意不把它归于人名下**。这条发现对本文最初的承重物（以时刻归属为唯一维度的”他择存量”）是不利的——按那个版本，民航应当已经拥有毕业生所缺的东西，而它拥有却不用。本文据此把承重物改到表 1b 最后一行：民航持有那一格的两个条件而缺第三个（归己），与反兴奋剂缺”无事入账”、信贷缺”他择”并列。**这是一次读取结果之后的改判**，按预注册纪律须明写：审计二的三条预测在改判前后都成立，改的不是预测，是预测所服务的命题；改判之后的命题在 7.2 加了第 7 条证伪条件，并在 11.3 加了第三条预测。

代理变量与限制：一份咨询通告是监管者的建议性文件，不等于航空公司实际做法；P1 的”定期评估”是定时他择，与本文核心那一格仍差一档。

7.6 两次审计的合读

教育侧与民航侧停在同一格（他择·定时），两侧的推行者都已自述该格在失效（贝尔法斯特的可穿戴设备；FAA 引用的事故报告）；两侧都没有进入无预告档。差别在：民航持有无预告档的存量而不归人，教育侧连存量也没有。这使 7.2 第 2 条的次序命题得到两侧一致的支持，同时把第 6 条存量条件的难点从”有没有”改成了”归不归”。

7.7 机制条件的完整检验设计

7.2 第 1 条是本文最贵的一条，两次审计都没有碰它。一个可执行的设计如下，供后续研究或愿意开放数据的机构使用。

- **单元**：同一机构内、同一年入职、同一职类的毕业生， $N \geq 30$ ；不跨机构、不跨国家，以避免在文化、规模、制度上处处不同的异质对比。
- **自变量**：按 5.6 编码规则对每人入职前的全部评价材料计算他择率，并分记无预告档。
- **控制变量**：学业成绩、作品集评分、结构化面试分、院校层级。
- **因变量**：入职后十二个月的绩效评价；若机构有“临时任务处理”记录，另取该项。
- **预测**：控制上述变量后，无预告档份数与十二个月绩效之间存在正向剂量-反应关系；他择·定时档的关系弱于无预告档；自择项数与绩效无关或负相关。
- **判负**：无预告档系数不显著或为负，本命题为伪；若他择·定时档系数不低于无预告档，5.2 第一根轴的三档划分为伪，退回两档。
- **样本纪律**：同质多单元优先于异质双单元；不得以两个国家或两个行业充数。

这个设计的最便宜版本只需要一家愿意回看自己招聘档案的机构。本文作者未取得这样的数据，故把它写成设计而不是结果。

7.8 审计三：那一格真的是空的吗

- **预注册（写于搜索之前）**：在四个有公开制度文本的职业场域各查一类记录——美国注册会计师协会的事务所同行检查；英国医学总会的医师执业再认证；英国照护质量委员会对照护机构的无预告检查；美国反兴奋剂机构公开的运动员检测历史。对每一类按表 1b 的三个条件逐一判：时刻归谁定/取样之手归谁/记录是否归己且记无事。判负条款：任一类三条件同时满足且记录可由被读者携带流通 → 7.2 第 7 条触发，承重物退回”他择存量”。停止规则：四类判完即停，不因结果补样。
- **结果（表 4）**：

表 4 四个职业场域的反例搜索

场域	时刻	手	记录归谁、记什么	是否填满那一格
会计师事务所同行检查 (AICPA)	他择·定时 (三年一轮，日期事先安排)	他手 (外部检查员)	归事务所而非个人；记”通过/有缺陷/未通过”	否：差他择·无预告与归人两项
医师执业再认证 (GMC revalidation)	自择 (五年一轮，年度考评材料由医师自行收集)	自手 (支持材料由本人整理)	归己	否：这是表 1b 第一行的职业版
照护机构无预告检查 (CQC)	他择·无预告	他手	归机构，公开评级；不归任何个人	否：与卫生等级牌同格
运动员检测历史 (USADA 公开页)	他择 (含无预告)	他手	按运动员姓名公开每季度检测次数；不记”结果无异常”，只在有处罚时另列；由机构持有，本人不能携带	否，但最近：三个条件占了两个半——差”记无事”与”本人持有”各半

判负条款未触发。但第四行必须单独写：美国反兴奋剂机构从 2010 年代起按姓名公开每名运动员每季度被检测的次数，运动员在争议中常引用”我被测过几十次”作为清白证据

——这是四个场域里唯一把他择·他手的次数归到个人名下并公开的做法。它没有记“无事”（次数不等于结果），也不由本人持有（页面归机构，不能带到另一个联合会），所以那一格仍是空的；但它说明这一格的填法已经有一半在实践里，且第一个走到这里的，仍然是三家里最早拿走时刻的那一家。7.2 第 7 条继续有效，11.3 预测三的“不算命中”清单据此加一句：仅公开次数不记结果的，不算。

限制：四个场域是作者按“有公开制度文本”选的，不是穷举；每类只读制度文本与公开页面，未查个案。

8 与最近的既有说法的分界

本节对 2 节回顾的各支逐一给出分界与判定形式；每一位都要说清它在自己领域回答的是哪个问题、为什么它没有吸收本文命题、若它成立而本文不成立会看到什么。

信号理论与筛选理论。它们要求信号的制造成本与能力负相关，而生成式 AI 使自择产物的制造成本对所有人趋于零。判定形式：若信号理论成立而本文不成立，招聘方应转向更贵的自择产物；若本文成立，应转向不由候选人选时刻的读数，哪怕很便宜。两条预测相反。残障原理同此，其昂贵落在信号本身，本文的可信落在时刻归属。Caplan (2018) 不区分信号的时刻归属；若其成立而本文不成立，AI 之后学历的信号价值应不变（成本未变）；若本文成立，学历应随其他自择产物一起下跌。

前台与后台 (Goffman, 1959)。后台是空间，且一旦被观察即成前台；本文的他择是时刻归属，不要求看后台——航线检查在前台做，检查员就坐在旁边，它之所以他择，只因这段航线不是飞行员选的。判定形式：若 Goffman 成立，无预告但公开在场的观察应与预告的观察等值；若本文成立，两者不等值，差别可从随后失效率读出。

自动化的反讽 (Bainbridge, 1983)。其结论落在保留技能；本文说残余可造、且残余记的不是缺的那部分，结论落在从哪些时刻读。Casner 等 (2014) 站在本文一侧。

工作样本测验 (Schmidt & Hunter, 1998)。效度最高的预测指标之一，但它是定时、预告、标准化的——他择·定时档，候选人会为它准备。本文预测 AI 之后标准化工作样本效度下跌而无预告成分不跌；可在下一轮元分析中证伪。

检查博弈。把不预告当检查者的策略；本文把它当证据的属性并主张累积为被检查者的存量——检查博弈中不存在“被检查者从被检查而无事中获得资产”这一步。判定形式：检查博弈预测检查频率与违规率的关系；本文预测他手无事账与随后被信任程度的关系，后者在检查博弈中没有变量。混沌工程读的是系统，本文读的是人。

卫生等级牌 (Jin & Leslie, 2003)。离本文最近的正面占位者：它已把无预告检查结果累积并公开展示，且证明该存量改变行为。本文未被吸收之处在表 1b 里是一行之差：等级牌落在他择·他手·不归己，本文的承重物在它下面一行。它证明了那一格的三个条件中两个可以同时成立并改变行为；本文说的是第三个条件成立时会出现一类新记录。若第三个条件在任何场域被证明不可成立，本文退化为等级牌制度在人身上的推广建议——7.2 第 7 条正是为此设的。

评估安全 (Dawson, 2021)。它管取样的条件归谁（环境、工具、身份），本文管取样的时刻归谁；两者在“预告的监考”格分开——条件全归校方而时刻公布，按评估安全已安全，按本文仍是他择·定时、可准备。判定形式：若 Dawson 成立，条件受控的预告考试效度在可穿戴设备普及后应不变；若本文成立，它应下跌而无预告成分不跌。

口试与同步核验层 (Joughin, 2010 ; Sotiriadou et al., 2020) 。已在做评估者当场定内容。本文未被吸收之处两处：它把口试当核验作品作者身份的手段（作品仍是证据，口试验真），本文把作品移出证据栏、口试本身成为证据；它不累积——口试无事不留痕，落在表 2 右列。判定形式：若口试支成立，口试成绩应与作品成绩高度相关；若本文成立，两者应在 AI 之后逐年脱钩，幅度可从同一门课的两个分数序列读出。

双车道评估。分界已在 7.4：它定义的是他择·定时格，本文的核心在第三档。

试用期。它是用人单位持有的他择·定时读数：时刻由单位定、被雇者不能推迟、期限公布。本文未被它吸收之处在于两点：试用期的结果不可携带（只对该单位有效，不进入任何被雇者可持有的记录），且它以“通过／不通过”记账，无事不累积。判定形式：若试用期已足够，用人单位在 AI 之后应延长试用期而非改招聘环节；若本文成立，用人单位应把可准备的招聘环节替换为现场环节，并开始把现场环节的结果写进候选人可持有的记录。

过程性证据。AI 时代评估文献的另一支主张以草稿、版本记录、反思日志等“过程性证据”替代终稿 (Boud & Molloy, 2013 的延伸；Frontiers in AI, 2026) 。这一支与本文同样不把证据落在产物上，但过程记录仍是自择的——学生决定何时保存一个版本、写一条反思；该支文献自己也承认过程材料“不能免于伪造”。按表 2，过程性证据落在自择·累积那一格，与作品集同格，只是更厚。判定形式：若过程性证据支成立，要求提交版本记录应当足以恢复评估效度；若本文成立，版本记录的效度将随生成工具对“生成过程”本身的模拟能力而下跌。

同批相邻研究 (本站) 。《能力证据债》把债记在证据的数量上，本文记在时刻归属上——补更多自择证据不减此债。《收方工序》答的是怎样给产物打折，本文答的是不折产物、换读数，两者可并存。《终稿之前》把风险记在路径上，本文记在时刻上——手写草稿仍是自择证据。《空签》《弃项》找的是被删除的分支，本文找的是被读取的时刻的归属，无正面重叠。《快，是因为那里已经判过了》解释自择产物为何变快变好，是本文目标格的成因。

9 伦理与适用边界

9.1 读数自带的三条条款

他择率、归己率与他手无事账几乎必然会被拿去考核人，故自带三条不可删的条款：

- (1) 它们指的是**证据的位置**，不是人的价值——他择率为零的人只是没被在他择时刻读过；
- (2) **不得**为了把这个数做好看而在任何安全关键岗位安排无预告取样——反兴奋剂的行踪规则有对价、有自愿、有范围，医院、机场、电网里的不预告干扰不是取样而是事故；
- (3) 已按此读数做出不利安排的机构，**必须公开其判定**“取样时刻由谁定”的编码规则并给出复核通道——一次预告三天的考核算哪一档必须写在规则里。

9.2 不得执行的场合与补偿

不得执行：未成年人；任何未自愿进入“检查库”的人；被读者无法拒绝进入的关系（雇佣中的下级对上级）。执行后的补偿：每一次他择读数，无论结果，都应写进被读者可持有的记录——否则它只是监视。

9.3 反噬预言

本命题一旦被制度采用，最省事的实现方式有两条，两条都会摧毁它要保护的东西。其一，把“无预告”做成一个可申报的项目——机构宣布自己有无预告制度，而实际时刻通过内部渠道被提前知道；行踪规则的运行史中这正是被反复揭出的漏洞。其二，把他手无事账

做成一张证书——一旦它成为可放进简历的纸，被评价者选择何时去考取它，第一个条件（他择）就被第三个条件（归己）吃掉了：归己是这一格最后一个条件，也是最容易反过来毁掉前两个条件的一个。信贷的历史已经演过一次——它最早把无事归己，也最早出现制造无事的市場。作者看不到出口。

9.4 不适用的领域

没有任何第三方愿意或能够不预告地读你的领域——纯个体创作、自由职业早期、以及 6.1 所说的所有未进检查库的人；成果本身不可能被外力制造的领域——若有一种产物 AI 造不出，那里自择证据仍有效，本文不需要，作者未找到这样的领域但不排除其存在。

9.5 与法律的接口

他择读数在多数法域都触到两条法律边界。其一，对劳动者的不预告观察受劳动法与个人信息保护法约束：欧盟一般数据保护条例与中国《个人信息保护法》都要求对个人的监测有明确目的、告知与最小必要；行踪规则之所以能运行，靠的是运动员签署的进入检查库协议，这一协议在教育与雇佣场景没有现成对应物。其二，把他择读数写入被读者可持有的记录，涉及记录的归属与更正权：信贷档案有《公平信用报告法》式的争议与更正程序，学业记录有学生隐私法的调阅限制，而“他择存量”作为一种新记录，其争议程序尚不存在。9.1 第三条条款（公开编码规则并给出复核通道）是对第二条边界的最低回应，不是充分回应。本文把这两条边界当作命题的适用前提，而不是待克服的障碍：没有协议与更正程序的他择读数，按 9.2 不得执行。

9.6 对三类行动者的含义

本文不给处方，但命题若成立，三类行动者各有一件事会变。

大学。成绩单上现有的信息是课程名、学分、等第；按本文，它缺一个字段——每一项评估的取样时刻归属。双车道大学已经在内部区分安全与开放两种评估，把这个区分印到成绩单上是 11.3 预测一押的那一步——它给的是第三个条件（归己）。做了这一步，成绩单就从一份自择证据的清单变成一份混合清单，用人单位第一次能从成绩单上读出他择率。但成绩单只记分数不记“无事”，且安全车道是定时的：即使做了这一步，大学仍停在表 1b 倒数第二行；要进最后一行，还差无预告成分与无事入账，那是预测二与预测三押的。

用人单位。校园招聘流程里可准备的环节（简历、作品集、可预约的面试、公布题型的笔试）在本文命题下效度趋零；能替代它们的不是更难的同类环节，而是候选人不能选时刻的环节——当场给一个真问题、给一段真材料、给一次真的打断。这类环节便宜，但要成为证据必须被记录并留给候选人：一次现场任务无论结果如何都写进候选人可持有的记录，它才是他择存量而不只是一次筛选。

毕业生。自择证据仍然需要——它决定你能否进门；但进门之后能积累的东西只有一种：主动进入有人会不预告地读你的位置。6.2 已说明这首先是机会的分布问题；对个体而言，可做的选择是在同等条件下优先选择带抽查制度的岗位与项目，并要求把那些时刻写进自己的记录。这不是一条励志建议，而是一条会计建议：他手无事账只在三个条件同时被记下的地方存在，而三个条件里毕业生自己能争取的只有最后一个。

10 局限与自反

10.1 方法限制

三个参照案例是作者选的，选源本身可能带有偏向；三次审计样本极小（六页、一份文件、四类制度文本），只能证明“未被推翻”，不能证明成立；审计一限于英语世界；机制条件（7.2 第 1 条）所需数据本文未取得。5.4 的他择率尚无任何一批真实毕业生样本被计算过。

10.2 本文解释不了的两个洞

其一，为什么反兴奋剂在护照失窃后没有像信贷那样出现公开的“制造他择读数”市场？是因为取样由第三方执行吗？若是，教育侧的他择读数一旦由校方执行，会不会立刻出现替考市场？审计二把这个洞挖开了一半：民航与反兴奋剂的差别被写进了表 1b 的第三条归属。没有挖开的另一半是：**为什么归己在民航被明令禁止**——安全文化的理由是归人会杀死报告；若这个理由对教育与雇佣同样成立（学生与雇员一旦知道无事也入账，会不会为了账而改变在他择时刻的行为），那么他手无事账在设立之日就会开始改变它要读的东西。9.3 把这写成反噬，但本文没有答案。其二，他择存量与“未见异常”的关系：一次他择读数出了异常，应从存量扣除还是另立一栏？三家做法不同（反兴奋剂只记异常，信贷两者都记），本文未裁定。

10.3 与本文命题相抵的已知事实

有一件事实与本文的方向相抵，必须写下：他择取样本身的抓获率极低。世界反兴奋剂机构 2012 年的工作组报告承认，尽管检测数量大幅增加，自 1985 年以来阳性率没有统计学上的改善；挪威反兴奋剂机构 2003–2019 年近五万份样本的违规检出率为 0.44%（Lauritzen, 2023），而以随机应答法估计的用药流行率远高于此。也就是说，即使证据落到了他择时刻，读出来的东西仍然很少。这对本文有两种读法。不利的读法：他择读数并不比自择读数更能分辨，本文的命题只是把一种低效换成另一种。有利的读法：抓获率低说明的是**只记异常**这一记法的失效——三家里唯一累积无事的信贷，恰恰是唯一不靠抓获率运转的一家；他择存量的价值在累积无事，不在抓到异常。本文倾向第二种读法，但承认第一种没有被排除；7.7 的设计可以分开两者——若无预告档的份数与绩效相关，而无预告档中的异常次数与绩效不相关，则第二种读法成立。

10.4 自反

用本文自己的判据给本文定位：这篇论文是一份**自择证据**。写作时刻是作者选的，提交时是准备好的，他择率为零。它能获得的唯一他择读数是同行评审——评审者在作者不选择的时刻、用作者不选择的标准读它。按本文命题，这份稿子的价值不在它自己写得怎样，而在它经不起那一次它不能自己安排的读取；它不能靠再修改一遍提高可信度——修改是自择的。这使本文自己落在表 2 的目标格里，与它批评的作品集同格。再进一步：本文用来推翻共有前提的材料（微量促红素不被护照标出）是一份他择读数——那个实验的受试者不选择取样时刻。本文最承重的一步靠的正是它主张的那类证据的两个条件（他择、他手）；这不是自封，它意味着若他择·他手的读数不比自择读数更可信，4.3 也就不比三家的立场更可信。而第三个条件本文自己不满足：这篇稿子若通过评审，评审记录归期刊持有，不归作者——本文自己也落在表 1b 倒数第二行。

11 结论与可检验预测

11.1 结论

生成式 AI 时代，大学毕业生的核心竞争力不是他交出来的东西的质量，不是机器替不了的那一小段技能，也不是他被记录得有多厚——三处都可被制造，三家各自的实证材料已经证明。剩下的是三条此前无人当作对象的归属——时刻归谁定、取样之手归谁、记录归谁持有——以及它们张成的八格里那一格四方皆空的位置。竞争力将是落进那一格的记录累积起来的量：由他人之手、在本人不选择的时刻取得、结果为无异常、且记在本人可以带走的账上。四个场域各自持有它的两个条件而无一持有第三个；三家各花了一代人走到自己那两个条件，教育侧可以不必再花一代人。

11.2 若命题为伪，会学到什么

若 7.2 第 1 条判负——无预告读数与绩效无关——本文会因此学到一件具体的事：分辨力不来自时刻归属，而来自别处，最可能的候选是任务的真实性（真问题与练习题的差别）而非时刻。届时应当检验的是“真任务”轴而非“他择”轴，而三家的参照案例仍然有效，只是它们各自抓住的东西须重新命名。若 7.2 第 2 条判负——有大学一开始就以无预告取样定义安全车道——本文会学到教育侧不必重走反兴奋剂的一代人，次序可以跳，那时应当研究的是为什么这一家能跳而另三家不能。若 11.3 预测三到期未命中，本文学到的是：那一格的第三个条件（归己）在人身上不可成立，民航的去身份与反兴奋剂的只记异常是常态而非例外——那正是 10.2 第二个半个洞的答案，届时本文的承重物退回表 1b 倒数第二行。

11.3 写死日期的两条预测

预测一：到 2028 年 12 月 31 日，澳大利亚八大集团与英国罗素集团合计三十二所大学中，至少三所公开说明学生的正式成绩记录或成绩单上区分“安全车道”与“开放车道”的评估项目。不算命中：仅在校内学习管理系统用图标标注、不进正式记录的；仅在政策文件承诺而未在记录上实施的；由第三方平台而非校方出具的。若到期未命中，5.5”这一栏将出现”的预测为伪，本文退化为一评估设计建议。

预测二：到 2029 年 12 月 31 日，上述三十二所大学中，至少五所在校级安全车道定义里加入一项由评估者当场决定内容或时刻的成分（现场追问、即时口试、随机抽取已提交作品的一段要求当场重做）。不算命中：预告题库的口试；由学生选时段的在线监考；单门课程或单个院系已在做的口试（8 节已说明其存在）——必须写进校级安全车道定义。若命中，教育侧正在走反兴奋剂走过的路；若到期不命中而安全车道仍被宣布有效，7.1 的侵蚀预测为伪。

预测三（改判后新增）：到 2029 年 12 月 31 日，在任一公开可查的场域出现一类记录，同时满足：取样时刻不由被读者定、由第三方执行、结果为无异常时也入账、且被读者可以调阅并携带到另一机构。不算命中：仅机构内部持有的；只记异常的；被读者可以选择时段的；仅公开检测次数而不记结果的（7.8）。若命中，表 1b 最后一行从此不空，本文的命题得到第一个实例；若到期未命中，7.2 第 7 条不触发，但 5.5”这一格将被填上”的判断降为未决。

参考文献

- 0a. American Institute of CPAs (AICPA). *Peer Review Program Standards* (现行版). 0b. Care Quality Commission (CQC). *How we inspect and regulate* (现行版). 0c. General Medical Council (GMC). *Guidance on supporting

information for appraisal and revalidation* (现行版) . 0d. U.S. Anti-Doping Agency (USADA). *Athlete Test History* (公开数据库, 按季度更新) . 1. Arrow, K. J. (1973). Higher education as a filter. *Journal of Public Economics*, 2(3), 193–216. 2. Ashenden, M., Gough, C. E., Garnham, A., Gore, C. J., & Sharpe, K. (2011). Current markers of the Athlete Blood Passport do not flag microdose EPO doping. *European Journal of Applied Physiology*, 111, 2307–2314. 3. Avenhaus, R., von Stengel, B., & Zamir, S. (2002). Inspection games. In *Handbook of Game Theory with Economic Applications* (Vol. 3, pp. 1947–1987). Elsevier. 4. Bainbridge, L. (1983). Ironies of automation. *Automatica*, 19(6), 775–779. 5. Basiri, A., Behnam, N., de Rooij, R., Hochstein, L., Kosewski, L., Reynolds, J., & Rosenthal, C. (2016). Chaos engineering. *IEEE Software*, 33(3), 35–41. 5a. Boud, D., & Molloy, E. (2013). Rethinking models of feedback for learning. *Assessment & Evaluation in Higher Education*, 38(6), 698–712. 6. Bureau d’Enquêtes et d’Analyses (BEA). (2012). *Final Report on the accident on 1st June 2009 to the Airbus A330-203, flight AF 447*. Paris. 7. Bridgeman, A., & Liu, D. (2024, July 2). Frequently asked questions about the two-lane approach to assessment in the age of AI. *Teaching@Sydney*. 8. Caplan, B. (2018). *The Case Against Education*. Princeton University Press. 9. Casner, S. M., Geven, R. W., Recker, M. P., & Schooler, J. W. (2014). The retention of manual flying skills in the automated cockpit. *Human Factors*, 56(8), 1506–1516. 10. Consumer Financial Protection Bureau (CFPB). (2015). *Data Point: Credit Invisibles*. Washington, DC. 11. Corbin, T., Dawson, P., & Liu, D. (2025). Talk is cheap: Why structural assessment changes are needed for a time of GenAI. *Assessment & Evaluation in Higher Education*, 1–11. 12. Curtis, G. J. (2025). The two-lane road to hell is paved with good intentions. *Higher Education Research & Development*, 44(8), 2151–2158. 13. Dawson, P. (2021). *Defending Assessment Security in a Digital World*. Routledge. 14. Dresher, M. (1962). *A Sampling Inspection Problem in Arms Control Agreements: A Game-Theoretic Analysis* (RM-2972). RAND. 14a. FAA Performance-based operations Aviation Rulemaking Committee / Commercial Aviation Safety Team (PARC/CAST). (2013). *Operational Use of Flight Path Management Systems*. Washington, DC. 15. Federal Aviation Administration (FAA). (2013). *SAFO 13002: Manual Flight Operations* ; (2017). *SAFO 17007: Manual Flight Operations Proficiency* ; (2022). *Advisory Circular 120-123: Flightpath Management*. 16. Goffman, E. (1959). *The Presentation of Self in Everyday Life*. Doubleday. 17. Jin, G. Z., & Leslie, P. (2003). The effect of information on product quality: Evidence from restaurant hygiene grade cards. *Quarterly Journal of Economics*, 118(2), 409–451. 18. Joughin, G. (2010). *A Short Guide to Oral Assessment*. Leeds Metropolitan University / University of Wollongong. 19. Kofinas, A., et al. (2025). Can graders distinguish AI-generated from student-authored work in authentic assessment? *Assessment & Evaluation in Higher Education* (在线先发) . 20. Lauritzen, F. (2023). Intelligence-based doping control planning improves testing effectiveness. *Drug Testing and Analysis*, 15(4). 21. National Transportation

Safety Board (NTSB). (2014). *Descent Below Visual Glidepath and Impact With Seawall, Asiana Airlines Flight 214* (AAR-14/01). Washington, DC. 22. OECD. (2019). *OECD Future of Education and Skills 2030*. Paris. 23. Schmidt, F. L., & Hunter, J. E. (1998). The validity and utility of selection methods in personnel psychology. *Psychological Bulletin*, 124(2), 262–274. 24. Sotiriadou, P., Logan, D., Daly, A., & Guest, R. (2020). The role of authentic assessment to preserve academic integrity and promote skill development and employability. *Studies in Higher Education*, 45(11), 2132–2148. 25. Spence, M. (1973). Job market signaling. *Quarterly Journal of Economics*, 87(3), 355–374. 26. World Anti-Doping Agency (WADA). (2009 及后续版本). *Athlete Biological Passport Operating Guidelines*; *International Standard for Testing and Investigations*. 27. World Economic Forum. (2020). *The Future of Jobs Report 2020*. Geneva. 28. Zahavi, A. (1975). Mate selection—a selection for a handicap. *Journal of Theoretical Biology*, 53(1), 205–214. 29. 站内相邻研究：孔凡鹤《能力证据债》（AI 时代九篇）；王德生《收方工序》（2026-08-29）；《终稿之前》（2026-09-01）；《空签》（2026-08-20）；《弃项》（2026-09-01）；《快，是因为那里已经判过了》。均见 sdeuniverses.com。

附录 A 两次审计的预注册原文

审计一（写于 2026-09-02，读取页面之前）：检索词 “two-lane approach assessment secured lane open lane universities”；取前六个大学官方或准官方页面；字段 (a) 监考／到场、(b) 校方定时、(c) 无预告取样；判负：≥2 所以 (c) 定义安全车道 → 次序命题撤回；不利结果照抄；六所读完即停。

审计二（写于 2026-09-02，打开 AC 120-123 之前）：P1 出现”定期评估手动飞行熟练度”或同义表述（预期出现）；P2 FOQA／LOSA 被列为数据源（预期出现）；P3 出现 “unannounced”或 “no notice”（预期不出现）；判负：P1 不出现 → 民航侧命题为伪；P3 出现 → 次序命题为伪；一份文件读完即停。

审计三（写于 2026-09-02，搜索之前）：四类记录——AICPA 事务所同行检查、GMC 医师再认证、CQC 无预告检查、USADA 运动员检测历史；每类按时刻／手／账三条件判；判负：任一类三条件同时满足且可携带 → 7.2 第 7 条触发；四类判完即停。

附录 B 审计一：六所大学安全车道定义的字段值

页面	(a) 监考／到场	(b) 校方定时	(c) 无预告	备注
悉尼大学（转述页）	是	是	否	仅允许能通过监考控制 AI 的评估类型；开放车道默认允许 AI
澳大利亚天主教大学	是	是	否	安全／开放两车道
奥克兰大学	是	是	否	AI 仅在”受控评估”可被限制；2025-11 修订程序
巴斯大学（两页）	是	是	否	2026/27 起以封闭／开放取代三色分级

贝尔法斯特女王大学	是	是	否	自述可穿戴设备使监考越来越困难
-----------	---	---	---	-----------------

附录 C 审计二：AC 120-123 全文字段计数

检索项	命中次数	读法
“manual flight”	24	手动飞行仍是主题之一
“proficiency”	38	熟练度被反复要求
“evaluat*”	23	含” 应被定期评估其手动飞行熟练度”
“instructor”	20	教员／检查员考虑事项五处单列
“LOSA” / “FOQA”	3 / 3	均作为训练与政策反馈数据源，无一处作为个人证据
“line check”	0	航线检查未以此名出现
“unannounced” / “no notice”	0 / 0	P3 不出现

计数由 pdftotext 转换后不区分大小写检索得出；全文 2,112 行。

声明

利益冲突：作者声明无利益冲突。

资助：本研究未获得任何资助。

数据可得性：两次审计所用材料均为公开网页与公开 PDF，检索词、日期与字段规则见附录 A-C，可复算。

伦理：本研究不涉及人类受试者或个人数据。

人工智能使用声明：题目、题型与研究要求由作者给出；三个参照案例的选取、共有前提的提出与推翻、命题、辨别装置、读数、证伪条件、两次预注册审计的设计与执行、以及全文写作，由作者与大语言模型 Claude (Anthropic) 协作完成；审计数据取自公开网页，检索词与样本规则在读取结果之前写死。作者对全文内容负责。

版本：第 3.0 版（终稿），2026 年 9 月 2 日。正文约 2.2 万汉字。沿革：1.0（1.34 万字）→1.1（补两位应用领域正主、辨别格改 2×3、增民航侧审计）→2.0（论文体、编码规则、检验设计、相抵事实）→2.1（承重物由”他择存量”改为”他手无事账”，改判来由见 7.5）→3.0（增审计三反例搜索，统稿）。

前五部把发生差从本体一路押送到可信的读数，最后一步是所有制度问题的最后一步：**记在谁的账上。**《入账之前的技能》守着这一层，材料是一百余年的商业秘密判例——法院是社会中最唯一被制度化地要求回答”这项技能属于谁”的机构。

它带来两条其余九章都不供给的东西。第一，**入账的本体论：**把技能写入任何第三方可复读的载体，不是记录，是触发一次不可回复的归类判定，且结果不由持有者控制——发生差的单向门在记录侧的镜像。第二，**载体的二分：**账本有记”如何做”的文件载体与记”为谁做”的名单载体，归属规则不同；AI 能吞下全部前者，进不了后者——它不是任何一份在后者名单上的一方。于是全书的制度落点收束在一句话上：技能保值的争夺，已从”能否编码”整体转移到”名单记在谁的账上”。

这一层同时是全书处方最锋利也最危险的地方：先于他人入账、只入重复项、让名单首记于自己账上——三条都可执行，三条也都有反噬。反噬的合流留给第十七章。

第六部 归属：账本层

第十四章 入账之前的技能

【章首】本章给方程收账：**归属层**。技能的价值不在技能之中，而在它尚未被写入任何账本的那一段——写入不是记录，是触发一次不可回复的归类判定，且结果不由持有者控制；账本有文件与名单两种载体，归属规则不同，AI 能吞下全部前者、进不了后者，于是保值的争夺整体转移到名单的记账权。判据一句：这项技能最近一次被写入账本，由谁写入？三十三件判例的预注册检验支持归类随入账形式而异，十一件客户关系案的探索显示名单归属随记账方而定、无一例外——两处不利结果（入账不等于归属转移；记忆组仍有受保护者）促成了本章两次自我修正，均照登。本章与本书其余九章互不指名，是审计乙判负的最大缺口所在，四十五对矩阵里凡涉本章的九条判定形式全部由第十五章补写并标〔补〕。本章欠的账：名单载体四十件判例的正式预注册检验（预注册已在正文自立），未跑；另欠同法域两部先行研究——员工知识归属的百年法律史与“人力资本法”一脉——各一段谱系交代，第十五章存目。

大白话：菜谱归店里，熟客名单跟着人走

一个厨师在一家店干了十年。他把拿手菜的做法写进了店里的操作手册。有一天他要走了，能带走什么？

菜谱带不走——它已经写进了店里的手册，从写下的那一刻起，它就是店里的东西了。

熟客带得走——那些人认的是他，不是这家店的招牌。他去哪儿，有一部分人跟着去哪儿。

《入账之前的技能》讲的就是这个差别，而且它的材料很硬：一百多年的商业秘密判例。法院是这个社会里唯一被制度化地要求回答“这项技能到底属于谁”的机构，所以这一章的证据不是理论推演，是判决。

两条要点。

第一条：入账不是记录，是一次不可回复的归类。

我们通常以为把东西写下来只是“留个记录”，东西还是自己的。判例说不是：**写下来这个动作本身，触发了一次关于“这属于谁”的判定，而且判定的结果不由持有者控制。**同样本事，你把它写进公司的文档，和它只留在你手上，法律地位是两回事。而且这件事不可回复——写下来之后，你不能通过“删掉”把它变回去。

第二条：账本有两种载体，规则不一样。

一种记“**怎么做**”——文件、手册、流程、代码、文档。另一种记“**为谁做过**”——客户名单、关系、谁欠你一次、谁认得你。

这两种载体的归属规则不同。而在今天，还有一个更要紧的差别：**AI 能吞下全部第一种，进不了第二种。**

它读得完世界上所有的手册，但它不是任何一份在场者名单上的一方——没有人跟它“认识”，没有客户因为它而回头，没有人欠它一次人情。

所以这一章的结论落得很实：**技能保值的争夺，已经从“能不能编码”整体转移到了“名单记在谁的账上”。**

读数是什么？

一句可以直接问的话：**这项技能最近一次被写入账本，是由谁写入的？**以及那个关键的追问：**你的名单，首记在谁的本子上？**

记在公司 CRM 里的客户，是公司的；记在你自己通讯录里、并且对方也认你的人，才跟着你走。这不是教人藏私，是让人看清楚一件本来就在发生的事。

一个小场景。

两个做了五年的顾问同时离职。甲这五年把所有方法论都沉淀进了公司知识库，业内评价很好。乙的沉淀少一些，但每个项目结束后都会亲自跟对接人吃顿饭，五年下来有几十个人认他。

甲离职时发现，他的方法论一句也带不走，而且公司立刻用它培训新人；乙离职时，有七八个老客户主动问他去哪儿。

这不是说甲做错了——沉淀方法论对公司有价值，对他自己的职业声誉也有价值。但**如果他以为那些沉淀是他自己的资产，他会在离职那天感到意外**。这一章的作用就是让这个意外提前发生。

它不是什么？

它不是“不要写下来”“不要分享”。写下来常常是对的，分享常常是对的。它说的是：**写下来的那一刻发生了一次归类，你至少要知道它发生了、结果是什么**。知情地交出去，和以为没交出去，是两回事。

再看两个身边的例子。

婚礼上的司仪。他手上有两样东西：一套流程稿，和几十对新人的家属还记得他。流程稿写下来就归了公司，第二天新人拿它也能照着走；而“某某家的婚礼是他主持的、办得体的面”这件事，记在几十个家庭的记忆里，谁也拿不走。两样东西的归属规则完全不同，而多数司仪只在意第一样。

老中医的方子。方子写下来就传出去了，写进教材就人人可查。但“这一片的老人他都认他、有病先找他”这件事，从来不在任何一本方书里。**方子是文件载体，认他的人是名单载体**。这个古老的区分在今天忽然重要起来，因为文件载体那一半正被整体吞掉。

三种常见的误读。

第一种：以为“名单”就是人脉。不完全是。人脉常指认识多少人；这里说的是**记账权**——那些人认你这件事，首先记在谁的本子上。一个人可能认识很多人，但那些关系全部是在公司系统里建立、由公司维护的，那么账不在他名下。

第二种：以为要少写、少沉淀。误读。写和沉淀本身常常是对的：它换来信任、换来合作、换来声誉，而声誉本身就是名单载体上最值钱的一笔。这一章要的不是少写，是**知情地写**——知道每一次写下都触发一次归类，而且结果不由你决定。

第三种：以为这是职场厚黑。不是。它的材料是判例，讲的是一件早就在发生的事：法律早就在裁定这些归属，只是多数人从没读过那些判决。**知道规则和钻规则的空子是两回事**——不知道规则的人，往往是唯一在离职那天感到意外的人。

给自己做个小检查。如果明天你离开现在的岗位，你能带走什么？逐样列出来，然后在每一样后面标一句：这样东西，首记在谁的账上？

它和旁边几个概念的区别。

和《现配优势》是同一现象的两种读法：一个从竞争差看（写进清单即被批量供给），一个从归属看（写进文件载体即进可复制区）。

和《零判份》接成一条因果链：入账→可复制→世界样本膨胀→判别力收窄，前面讲过。

和《单遍段》是一对镜像单向门：对象侧的门与记录侧的门，前面讲过。

和《常席差》有一个必须分清的地方：常席差说能力的读数不可裁归个人，而这一章的判例显示**关系名单可以按记账方裁归**。这不矛盾——**能力账和关系账是两本账**，各有各的规则。把客户名单当成能力读数来用，是实践中最常见的一种误读：名单证明的是关系的归属，不是这个人的能力有多强。

一句话记住它：这一章管的是**记在谁的账上**——而且账有两本，一本 AI 能吞下，一本它进不去。

怎么长出来：方法、技术与流程

个人层面。

第一，**知道你在往哪本账上写**。每次把一个做法写进某个系统之前，花三秒钟意识到：这一下是把它记进谁的账。写还是要写，但要知情。

第二，**别把名单托付给别人的系统**。你自己也留一份：谁、在什么事上、什么时候、我们之间发生过什么。这不是要你另起炉灶，只是让那本账不止一份，且有一份首记在你自己名下。

第三，**先自入重复项**。如果你有一样东西迟早要被写下来，那就自己先写一次并留下自己的版本，而不是等着它只以别人的形式入账。

大学发生学教育：教学生看懂账本，并给他一本自己的。

大学在这件事上有一个几乎无人注意的缺口：学生四年里所有的记录都记在学校账上——成绩记在教务系统，项目记在指导老师那里，实习记在企业那里。**毕业时他能带走的，只有一张成绩单和几张证书**。他四年里认识了谁、谁认他、他在谁那里留下过什么，全部没有账。

- **关系账本课（一次讲座就够）**。在低年级明确讲清两种载体的差别，以及入账即归类这件事。这是最省钱、也最容易被忽略的一课——很多人是在三十五岁离职那天才第一次明白它。
- **实习与合作项目的归属条款教学**。让学生在签任何实习协议、项目协议之前，读懂其中关于“你产出的东西归谁”的条款。不是教他去争，是教他**知情**。这一条应当进就业指导，而不是等他吃亏后再去补。
- **可携带的关系记录**。鼓励并帮助学生建立自己的一本账：本科四年里，哪些老师、企业导师、合作方、同侪，在什么具体的事情上与他共事过。学校可以提供模板与签字确认的便利，但这本账**必须首记在学生自己名下**，而不是只存在于学校系统里。
- **先自入重复项的训练**。学生在项目结束时，除了交给学校的版本，自己也留一份完整的复盘（含数据、思路、失败处）。这份自留版本是他自己账上的第一笔。

教师的动作。

在学生毕业前问一次这个问题："如果明天你离开这里，你能带走什么？"——大多数学生会答成绩单和证书。这个问题的价值不在答案，在于让他第一次意识到还有另一本账。另外，教师应当**主动成为学生名单上的一笔**：明确告诉学生"这件事上我认你，将来需要我说话可以来找我"。这句话对学生的实际价值，常常超过一封格式化的推荐信。

最容易做坏的地方。

把这一章读成"要藏私、要留一手"。那是彻底的误读，而且会毁掉一个人的合作声誉——合作声誉本身就是名单载体上最值钱的一笔。正确的读法是：**照常写、照常分享、照常沉淀，但知道每一次入账都是一次归类，并且给自己也留一本账。**

原文

入账之前的技能：人工智能时代技能保值的法律结构与实证检验

王德生

（作者简介：王德生，计算数学博士，德麦国际 Demai International Pte. Ltd. (新加坡) 首席执行官，曾任教于斯旺西大学、南洋理工大学。）

摘要：生成式人工智能显著缩短了技能从形成到被编码复制的时间，个人技能的贬值成为普遍现象。既有解释分属三个传统：商业秘密法主张价值来自信息的不为人知；表演研究主张价值来自不可复制的现场；技能习得研究主张价值来自持续重解问题的能力。本文指出，三种解释共享一个未经检验的前提，即凡使技能产生价值的要素均可被无损地记入某种账本。以美国商业秘密法一百余年的判例为材料，本文论证：员工的一般知识、技能与经验被排除在秘密保护之外这一稳定规则，表明存在一类要素，其一旦被写入账本，即触发一次不可回复的归类判定，且判定结果不由持有者控制。据此提出"前账阶段"概念，并区分账本的两载体——记录"如何做"的文件载体与记录"为谁做"的名单载体，主张人工智能只能吞并前者，因而技能保值的争夺已转移至名单载体的记账权。对 33 件公开判决的预注册检验显示：争议技能仅存于记忆的案件，法院判为不受保护一般技能的比例为 73%；技能以文件形式被带走的案件为 55%。探索性子样本显示，11 件以客户关系为争点的案件中，名单由员工自建或早于雇佣的 6 件全部判归员工，由雇主分配、记录或购入的 5 件全部判归雇主。文末给出可检验预测与政策含义。

关键词：技能贬值；商业秘密；一般知识技能经验例外；知识编码；客户关系归属；人工智能

Skill Before the Ledger: The Legal Structure of Skill Preservation in the Age of AI and an Empirical Test

Abstract: Generative AI has sharply shortened the interval between the formation of a skill and its codified replication, making individual skill depreciation pervasive. Three traditions explain skill value differently: trade secret law locates it in secrecy, performance studies in the unrepeatable live event, and motor learning in the capacity to re-solve problems. This article shows that all three share an unexamined premise—that every value-bearing element of a skill can be entered into a ledger without loss. Drawing on a century of U.S. trade secret doctrine, it argues that the stable exclusion of employees' general knowledge, skill and experience from secret protection reveals a class of elements whose entry into any ledger triggers an irreversible

classification decision outside the holder's control. The article names this class the "pre-ledger stretch" and distinguishes two ledger carriers—documents, which record how a skill is performed, and rosters, which record for whom—arguing that AI can ingest only the former, so that the contest over skill value has shifted to the entry right over rosters. A preregistered test of 33 public opinions finds courts classify memory-only skills as unprotectable general skill in 73% of cases versus 55% where documents were taken. An exploratory subsample of 11 customer-relationship cases finds rosters built by the employee or predating employment awarded to the employee in all six cases, and rosters assigned, recorded or purchased by the employer awarded to the employer in all five. Falsifiable predictions and policy implications follow.

Keywords: skill depreciation; trade secrets; general knowledge, skill and experience exclusion; knowledge codification; customer relationship ownership; artificial intelligence

一、问题的提出

一位从业十五年的数据库工程师在 2024 年发现, 他最擅长的慢查询诊断, 大型语言模型可以在三十秒内给出八成结论。他没有失业, 但报价随之下降。这不是孤例。Epoch AI (2024) 核对开放模型对封闭模型的复现时差: AlphaFold 为 35 个月, GPT-3 为 23 个月, DALL-E 为 15 个月; 从 2024 年 9 月的 o1 到 2025 年 1 月的 DeepSeek-R1, 只剩 4 个月。技能从形成到被复制的时间正在向零逼近。

人力资本理论对此的标准回应是区分可编码技能与不可编码技能 (Autor, 2014), 并预期价值向社交、判断与组织能力迁移 (Deming, 2017)。这一回应描述了迁移的方向, 但没有回答一个更基本的问题: **一项技能的价值究竟依附于什么, 以至于编码会使它消失?** 若价值依附于技能本身, 编码只是复制, 不应改变原件的价值; 若价值依附于技能之外的某种条件, 则需要说明那是什么条件, 以及编码为何摧毁它。

本文的主张可以概括为三句话。第一, 技能的价值不在技能之中, 而在技能尚未被写入任何账本的那一段——本文称之为"前账阶段"; 写入账本的动作不是记录, 而是触发一次不可回复的归类判定。第二, 账本有两种载体: 记录"如何做"的文件载体与记录"为谁做"的名单载体, 两者遵循不同的归属规则。第三, 人工智能改变的是文件载体的入账速度, 它无法进入名单载体; 因此在人工智能时代, 技能保值的核心不再是防止编码, 而是名单的记账权归属。

本文的材料来自美国商业秘密法与竞业限制法的判例。选择这一材料有两个理由。其一, 法院是社会中最唯一被制度化地要求回答"这项技能属于谁"的机构, 其判决构成关于技能归属的最系统的公开记录。其二, 商业秘密法内部存在一条延续百余年的规则——员工的一般知识、技能与经验不受秘密保护——而这条规则从未被用来回答本文的问题。

下文第二节梳理三个既有解释传统及其共享前提; 第三节提出前账阶段概念; 第四节区分账本的两载体; 第五节给出命题与可检验推论; 第六、七节报告预注册检验与探索性分析; 第八节讨论竞争解释、适用边界与政策反作用; 第九节结论。

二、三种既有解释及其共享前提

(一) 价值来自不为人知: 商业秘密法

美国《统一商业秘密法》第 1(4) 条将商业秘密定义为“因不为一般人所知且不能通过正当手段轻易获得而具有独立经济价值”的信息；1939 年《侵权法重述》第 757 条评注 b 已持相同立场。价值不在信息之中，而在“不为人知”这一条件之中：条件一旦丧失，价值归零且不可回复。Kewanee Oil Co. v. Bicron Corp., 416 U.S. 470 (1974) 进一步确认逆向工程与独立发现合法，因而秘密的价值仅存在于竞争者尚未获得的时段。

该领域内部关于保护依据存在长期争论。Bone (1998) 认为商业秘密法缺乏独立的正当性基础，寄生于合同法与侵权法之上；Lemley (2008) 主张将其视为知识产权的一种。两派的分歧在于法律为何保护“不为人知”，而非价值是否来自“不为人知”。

(二) 价值来自不可复制的现场：表演研究

表演研究给出方向相反的答案。Phelan (1993: 146) 主张表演唯一的生命在当下，它不能被保存、记录或以任何方式进入再现的流通，一旦被记录便成为他物。Auslander (1999) 批评这一本体论立场，指出“现场”范畴仅在“被中介”这一对立面出现后才具有意义，现场的价值是由录制条件生成的。Fischer-Lichte (2008) 从机制上补充：价值经表演者与观众共同在场的时间内被组织出来，作品对象只是其残余。两派的分歧在于现场的价值是先天的还是被建构的；两派的共识是：显现是价值发生的场所，不显现则无价值可言。这与商业秘密法的立场正面相反——一方视显现为价值毁灭的时刻，另一方视不显现为价值的缺席。

(三) 价值来自持续重解的能力：技能习得研究

运动技能习得研究提供第三种答案。Bernstein (1967) 提出“练习是一种不重复的重复”：熟练者并非重复某一动作，而是在约束每次变化的条件下重新解决同一运动问题。后续实验一再证实：熟练者的动作变异不减反增 (Latash, Scholz & Schöner, 2002)；变式练习在练习期表现较差而迁移更好 (Shea & Morgan, 1979)；技能是表演者、任务与环境三方约束下生成的关系属性，而非储存于个体的程序 (Newell, 1986；Davids, Button & Bennett, 2008)。该领域内部，Schmidt (1975) 的图式论主张个体储存抽象的运动图式，生态动力学则否认任何储存物的存在。两派的共识是：被记录下来的动作不是技能本身。

(四) 共享前提

三种解释分别将价值定位于条件（不为人知）、显现（现场）与过程（重解）。它们互相排斥，但共享一个未被言明的前提：**凡使技能产生价值的要素，都可以被无损地记入某种账本**——可以被指认为一个秘密、被看见为一次显现、被测量为一种变异。三者的分歧只在记账方法，“账本可以完整记录全部价值要素”这一点被三方共同视为不必论证。

上述前提可以进一步追溯至一个更一般的假定：技能的保值可以由某个终点单独结算——先将技能转化为一件可持有的终点物（一个秘密、一次事件、一个动作解），然后持有它。若保值可由终点结算，则一切价值要素自然可在该终点上被计入。

三、前账阶段：入账作为不可回复的归类判定

(一) 来自商业秘密法内部的反证

推翻上述前提的证据来自商业秘密法自身。自 1912 年起，美国各州法院持续拒绝将员工的“一般知识、技能与经验”计入受保护的商业秘密，即使这些技能对雇主而言确属秘密。CVD, Inc. v. Raytheon Co., 769 F.2d 842, 852 (1st Cir. 1985) 给出的理由被沿用至今：这一例外“实现了劳动力流动的公共利益，促进雇员从业的自由与竞争的自由”。Saunders 与 Golden (2018) 对判例的系统编码显示，法院主要考察四项因素：员工入职

前的经验、员工是否参与秘密的开发、秘密是否已被实际使用，以及离职时带走的是文件还是仅为记忆。Hrdy (2019) 进一步指出，法院系统性地漏判了一类信息——对公司而言技术上是秘密、却因长于员工身上而不受保护。

需要谨慎地限定这一材料所能支持的结论。法院从未表述为"入账使技能失效"；法院的表述是这些技能本来就属于员工，雇主的账本从未有权记载它们。本文不在判例之外多作推论，只指出判例中可以直接观察到的次序：**只要员工的技能未被写入任何账本，它就无需被归类；雇主将其写入保密协议、诉状或禁令申请的那一刻，法院被迫作出一次归类判定——判其为秘密，或判其为一般技能——而这一判定不可回复。**

Variable Annuity Life Insurance Co. v. Miller (Tex. App.—Amarillo 2001) 是最清晰的例证。同一名员工、同一项技能，书面记录与自计算机下载的部分被判归雇主并禁止使用，保留在员工记忆中的部分被判为一般技能不予禁止。入账没有改变技能本身，改变的是它必须被归类这一事实，以及归类一经作出即不可撤回。

这一观察否定了第二节（四）的前提：并非所有价值要素都可被平静地计入。存在一类要素，其入账即触发一次不可回复的判定，且判定结果不由持有者决定。

（二）概念界定

本文将这类要素命名为**前账阶段**（the pre-ledger stretch），并给出可操作的界定：一项技能中满足以下三个条件的部分——（1）尚未被写入任何第三方可复读的载体；（2）其价值的实现以特定在场者的存在为前提；（3）每次实现的内容随在场约束的变化而不同。三个条件同时成立，该部分即处于前账阶段；任一条件不成立，它就属于重复项、快照或已入账项。用一句话概括本文的核心判断：**保值的技能不是一个可持有的秘密，不是一次不可复制的事件，也不是一种可随身携带的能力，而是一段先于任何账本、且入账即触发不可回复归类的阶段。**

该概念的三个属性各有来源。**不可回复**来自商业秘密法：入账不可逆，公开之后不能再宣布为秘密。**在场**来自表演研究：耦合仅在有在场者时发生，无在场者的耦合是排练而非交付。**不重复**来自技能习得研究：耦合每次不同，任何记录都只是快照。

前账阶段不是三种既有解释中的任何一种。它不是秘密，因为它无需不为人知——即使被完整描述，学习者仍无法据此复制，因为下一次的约束已经改变。它不是事件，因为它可以反复发生，只是每次不同。它不是个体持有的能力，因为条件（2）与（3）使它无法脱离特定在场者与特定约束而存在——同一人在另一片约束下产生的是另一段耦合，同一片约束下换一人亦然。它是持有者与在场约束之间的关系，而非任何一方的属性；这正是技能习得研究所谓"关系属性"（Newell, 1986）在经济归属问题上的含义。

人工智能改变的正是这一阶段的长度。过去技能从发生到入账存在时差，价值在这一时差中被兑现；当前时差趋近于零，每次公开显现都可能构成入账。因此问题不再是"守住什么"，而是**谁来入账、入账什么、何时入账。**

四、账本两种载体

（一）归属矩阵

将一项技能的组成按入账前与入账后两个时点分类，可得下表。

表 1 技能组成的入账归属矩阵

入账前类别 \ 入账后类别	受保护的秘密	一般技能	公开项
---------------	--------	------	-----

重复项（每次相同、可录制、可教授）	文件被带走时可能落入此格（Dur-A-Flex v. Dy, Conn. 2024 ; Caudill Seed v. Jarro, 6th Cir. 2022）	多数案件落入此格	模型训练语料的来源
前账阶段（每次不同、仅在场约束下发生）	几乎无法落入此格：法院拒绝将员工身上的技能记为秘密	入账即触发归类，且多被判入此格	被录制后成为快照
在场者名单（为谁做）	记录于雇主客户系统者判归雇主（Hanger, Tenn. App. 2008 ; Cabela's, Del. Ch. 2018）	员工自建或早于雇佣者判归员工（VALIC, 2001 ; Combs, Tenn. App. 2013）	公开名录不归任何一方

矩阵中最值得关注的是第二行第二列：前账阶段一旦入账，即被迫归类，且多被记为一般技能。这一格有三个可观察的特征。其一，入账当年所有可读指标均在改善——简历更完整、保密协议更严密、证书更多、模型更能干。其二，被替换的一侧在两种记法中均无对应栏目——入账前它只被称为“会做”，入账后被称为“一般技能”，两本账均无“前账阶段”一栏。其三，使用即抹除——证书越多、文档越完整、录像越齐全，曾有耦合存在的证据越难查证，因为每一份证据本身都是一次入账。

（二）文件载体与名单载体

表 1 第三行揭示了一个此前未被区分的事实：账本有两种载体。**文件载体**记录“如何做”，包括文档、录像、代码、语料；**名单载体**记录“为谁做”，即在场者名单——客户、患者、听众、当事人。两种载体遵循不同的归属规则。对文件载体，法院考察带走的是文件还是记忆（Saunders & Golden, 2018 的第四因素）；对名单载体，法院考察名单由谁建立、记录于谁的账上（详见第七节）。

两种载体的入账者也不同。文件载体的入账者可以是任何人，包括机器；名单载体的入账者只能是账本持有人——一个人不能替另一个人建立关系。这正是人工智能所改变的边界：模型能够吞并全部文件载体，凡被写下的皆可进入；模型无法进入名单载体，它没有可以记录“谁为谁做”的账本，它不是任何一份在场者名单上的一方。因此，文件载体上的入账权已经失守——每次公开显现都是入账；**入账权的争夺整体转移至名单载体**。一项技能在人工智能时代能否保值，不取决于它能否被写下（一切皆可被写下），而取决于其在场者名单记录于谁的账上。

名单载体一格同样具有三个特征。其一，雇主将名单录入客户系统的当年，账面上是改善：客户关系被“资产化”，可在资产购买协议中转让（WorldVue Connect v. Szuch, 5th Cir. 2025）。其二，被替换的一侧在两种记法中均无栏目——员工建立名单这一事实，在雇主系统中称为“客户”，在员工简历中称为“业绩”，双方均无“由谁建立”一栏。其三，使用即抹除：客户系统使用越久、越规范，名单由员工建立的证据越难查证。Hanger Prosthetics v. Kitchens (Tenn. App. 2008) 中法院的措辞正是“在雇主的费用下发展起来的关系”，员工自建的那一段已无从证明。

五、命题与可检验推论

（一）三个命题

命题 1（不可复归类）：技能的价值要素一经写入任何第三方可复读的载体，即触发一次归类判定；判定不可回复，其结果不由持有者决定。

命题 2（两种载体）：账本存在文件载体与名单载体两种形式；文件载体的归属取决于载体形式（文件抑或记忆），名单载体的归属取决于记账方（由谁建立、记录于谁的账上）。

命题 3（载体转移）：机器能够进入的账本载体的边界，即机器能力的边界。生成式人工智能能够写入全部文件载体，不能写入任何名单载体——它不是任何一份在场者名单上的一方。因此在编码成本趋近于零的条件下，文件载体的入账权必然失守，技能保值的决定因素转移至名单载体的记账权。

（二）三阶段保值策略

由三个命题可推出一条具有操作性的保值路径，其阶段次序可由观察判定：若观察到工作方式先于呈现方式改变，即符合本文所述路径；若呈现方式先于工作方式改变，则属另一路径，本文预测其无法完成第三阶段。

第一阶段涉及工作方式：将"重复解"改为"在变化的约束下重解"，仅交付结果而不录入解法，并建立一份仅记录"本次约束与上次差异"的私账。这一阶段是持有者可以直接介入的环节，介入点是入账权——自行决定何时入账、入账哪一部分。

第二阶段涉及呈现方式：仅以在场事件的形式交付，将可录制的重复项与不可录制的耦合分开。Meyer-Dinkgräfe（2015）对英国国家剧院 2009 年以来现场直播的考察表明，录播未蚕食现场票房，付费者能够区分被录制的部分与在场的部分。

第三阶段涉及名单：积累一份持续要求重解的在场者名单，其付费对象是"下次仍需持有者在场"而非上次的产物。该名单必须记录于持有者自己的账上——由在场者主动发起，直接向持有者付费，其记录首次出现于持有者的账本。这是全部阶段中唯一机器与雇主均无法代行的环节。

路径的不可逆环节是入账本身：任何部分一经写入第三方可复读的载体，即永久进入可复制区。不入账并不可行——无入账即无报酬，证明持有技能需要痕迹，痕迹即入账。因此路径的要求不是不入账，而是**先于他人入账，且仅入重复项**。路径的失败点在第三阶段：在场者名单为零时，技能无论优劣均退化为重复项，因为无人为下一次付费。

（三）判据

上述命题可压缩为一个不依赖主观判断、可直接询问记录的判据：**这项技能最近一次被写入账本（保密协议、证书、录像、简历、训练语料），由谁写入——持有者本人还是他人？写入之后，下一次仍须重做的是哪一环节？**

表 2 判据在六类情境中的应用

情境	最近一次入账由谁写入	下一次仍须重做的环节
开源程序员	本人（提交至公开仓库）	新约束下的架构判断
外科医生	医院（手术录像入库）	每一病例的解剖差异
同声传译	客户（录音进入语料）	现场配对
连锁餐饮厨师	加盟商（写成标准作业程序）	当日食材调整——但已无人为此付费，名单归零
离职员工诉讼	雇主（写入禁令申请）	法院将"带走的是文件还是记忆"列为归类因素（Saunders & Golden, 2018）
剧场演员	剧院（现场直播）	录播未蚕食现场票房（Meyer-Dinkgräfe, 2015）

后两种情境并非由命题推出，而是由既有记录查证所得。第五种情境尤为重要：它表明本文的判据并非本文所创，法院一百余年来一直以入账形式判断技能归属，本文只是将这一判据从法庭移至个人的账本。

（四）竞争解释与可证伪条件

最直接的竞争解释有二。其一，默会知识说：Polanyi（1966）指出人所知者多于所能言者，Autor（2014）据此解释自动化在某些任务上的停滞。其二，互补资产说：Teece（1986）证明创新的利润常归于持有互补资产（制造、渠道、关系）的一方而非创新者。

本文与二者的差别落在一个它们无法预测的变量上：**入账次序**。默会知识说预测价值来自不可言说；本文预测前账阶段可被完整言说而仍无法被复制。互补资产说预测价值归于资产持有者；本文预测持有者自身亦无法携带前账阶段。两种竞争解释均预测：只要内容不可编码或资产在手，价值即存在，与谁先入账无关。本文的可证伪条件由此写为：**在控制技能的可编码性与雇主的互补资产之后，若"他人先入账"与"本人先入账"的技能在被复制的时点上不存在先后差异，则本文命题为伪**。届时应承认价值仅为默会程度或资产归属的函数。

最低成本的检验来自表 2 第五种情境：若法院以入账形式判定归类，则争议技能仅存于记忆的案件，被判为一般技能的比例应高于技能以文件形式被带走的案件。该检验所需数据在公开判决书中已经存在。

六、数据与方法

（一）预注册

在读取任何结果之前，本文写定以下事项。数据来源为 CourtListener 公开判例库，检索词为"general knowledge, skill"与"trade secret"的四种组合，取可获得全文的判决；样本量不少于 30 件。编码两项：争议技能是否以文件形式被带走；法院是否判为一般技能（不受保护）。判负条款：记忆组与文件组判为一般技能的比例差小于 15 个百分点，则命题 1 的经验推论判负。判负后的处置：将"法院以入账形式判归类"降为"法院部分以入账形式判归类"，删除第八节第四款的时间界预测。停止规则：编码满 30 件即停止。

（二）样本构成

四种检索组合共得 255 件唯一判决，其中 58 件可获得全文，39 件正文两次以上提及"一般知识／技能"。剔除与雇员无关的抢先适用案 1 件、文本损坏无法编码 1 件、重复上传 1 件、无法判定结论 3 件，得 33 件。其中 6 件为竞业限制案而非狭义商业秘密案，二者同样适用一般技能例外，但竞业限制案中"受保护"的含义为"雇主具有可保护利益"。

（三）编码规则

"以文件形式被带走"依判决书正文关键词初筛后逐件人工确认，"记忆组"包含法院未提及任何文件的案件。编码由单人完成，未经双人复核。样本来源于检索词而非全体判例的随机抽样。逐件编码见附表。

七、结果

（一）预注册检验

表 3 争议技能存在形式与法院归类

争议技能的存在形式	判为一般技能（不受保护）	判为受保护	一般技能比例
仅存于记忆（22 件）	16	6	73%

以文件形式被带走（11 件）	6	5	55%
----------------	---	---	-----

两组相差 18 个百分点，判负条款未被触发。最清晰的一件为 VALIC v. Miller (2001)：同一案件中，法院将书面信息、下载信息与软件判归雇主并禁止使用，将员工记忆中的信息判为一般技能不予禁止——同一人、同一技能，以入账形式一分为二。

（二）两项不利结果

第一，文件组仍有 55% 被判为一般技能。这表明入账不等于归属转移给入账者：雇主将其写入账本，法院仍可能判其为员工的一般技能。本文原拟表述为"入账即归属他人"，现修正为"入账即触发归类"：入账改变的是它必须被归类这一事实，而非它属于谁；归属由法院另行判定。

第二，记忆组中有 6 件被判为受保护，其中 4 件（Vantage Technology v. Cross, Tenn. App. 1999；Hanger, 2008；HCTec Partners v. Crawford, Tenn. App. 2022；WorldVue, 2025）并非基于文件，而是基于"特殊训练"与"客户关系"。WorldVue 一案中，雇主通过资产购买协议购入了"客户及客户关系"这一栏目。这表明账本存在第二种载体——关系名单——这正是第四节（二）的来源。

（三）探索性分析：名单载体的归属规则

将 33 件中以客户关系为争点的案件单独考察，共 11 件。此项分析不在预注册之内，系读取不利结果后进行，本文按探索性分析报告，并在第八节另立预注册。

表 4 名单载体的记账方与法院归属判定（探索性）

名单由谁建立、记录于谁的账上	案件	判归
员工自建，或早于雇佣（私人关系、公开名录、入职前客户）	VALIC v. Miller (2001)；Anderson Chemical v. Green (Tex. App. 2001)；Oak Mortgage v. Ameripro (Tex. App. 2015)；Total Staffing Solutions v. Staffing, Inc. (Ill. App. 2023)；Combs v. Brick Acquisition (2013)；Corbin v. Tom Lange Co. (Tenn. App. 2003)	6 件 全部判归员工
雇主分配、记录于雇主系统、或由雇主购入	Hanger (2008)；Vantage (1999)；Dill v. Continental Car Club (Tenn. App. 2013)；Cabela's v. Wellman (2018)；WorldVue (2025)	5 件 全部判归雇主

11 件无一例外。但须说明：该子样本系从为检验命题 1 而抽取的 33 件中事后划出，检索词并非围绕客户关系设计，且样本量小；它支持的是一个假设而非一项结论，正式检验依第八节（四）的预注册另行进行。就现有材料而言，法院判定名单归属所依据的不是名单上有谁，而是名单首次记录于谁的账上。Corbin 一案尤具说明性：公司刻意以"接力传话"的方式使每位客户同时认识多名销售人员，法院据此认定客户关系是公司的技术而非员工的关系——却仍判归员工，因为公司并未将其记于员工名下，反而稀释了它。名单不在任何一方的账上时，归于其实际发生的一方。

八、讨论

（一）与相关文献的分界

知识编码与模仿速度。Zander 与 Kogut (1995) 以瑞典制造企业数据证明知识的可编码性与可教性直接缩短其被竞争者模仿的时间，是本文第三节图景的实证原型。二者的差别仅在一列变量：他们可将编码性作为知识的属性测量，数据中不含"由谁编码、谁先编

码”；本文将入账视为有主体、有次序的行为。可判定的差别是：控制可编码性后，若本人先编码与他人先编码的技能在被模仿时间上无差异，则属性决定一切，本文命题 1 的经验含义落空。Cowan 与 Foray（1997）指出编码本身是一次创造，不完全替代其所依据的默会知识；本文同意，并补充编码同时是一次不可回复的归类。Nonaka（1994）将“外显化”视为组织知识创造的必经环节，与本文方向相反——前者关注入账创造的价值，本文关注入账触发的归类。可判定的差别是：外显化之后，若个人对该知识的议价能力上升，Nonaka 的框架成立；若个人下降而组织上升，本文成立。Winter（1987）的知识资产分类（可教／可表达／可观察）是这一文献族的源头，本文的“重复项”对应其中可表达且可观察的一格。

信息悖论。Arrow（1962）指出买方不见信息无法定价、见之则无需付费。前账阶段看似其一例。差别在于 Arrow 的对象是见之即可携走的信息；前账阶段见之不可携走，唯入账可携走。可判定的差别是：经保密协议中介的披露之后，若技能的价值恢复，则价值损失仅为交易结构问题；若不恢复，则因无可披露之物，披露的只是快照。

互补资产与人力资本专用性。Teece（1986）主张价值归于互补资产持有者；本文主张持有者自身亦无法携带前账阶段。可判定的差别是：员工携带其整个团队与设备跳槽时，若技能的价值随之迁移，则互补资产说成立。Becker（1964）的专用人力资本说预测离职即失值；本文预测不离职、仅入账亦失值——未离职员工将其技能写入公司知识库后，议价能力若不降，则专用性说成立。

业绩的可移植性。Groysberg（2010）追踪上千名跳槽的证券分析师，发现其业绩在换公司后普遍下滑，除非携团队一同离开，并据此认为业绩依附于公司资源。本文对同一数据的解读不同。Groysberg 的样本中，业绩保持的两类人——携团队离开者，以及跳槽至原雇主同类平台者——共有的并非公司资源，而是**名单的记账方未变**：前者带走了记录于自身或团队账上的客户，后者进入了同样以分析师个人为客户记账单位的机构。业绩下滑者则是把名单留在了原雇主的客户系统中，跳槽时带走的只有文件载体（研究方法、模型），而这部分恰是可复制的重复项。两种解读在一个可观察的变量上分手：Groysberg 的“公司资源”随公司而定，本文的“记账方”随客户由谁发起、向谁付费而定，二者可以在同一家公司内部不一致。可判定的差别是：控制团队是否随迁后，若名单记于自身账上（客户由本人发起、直接向本人付费）的分析师业绩亦随跳槽下滑，则 Groysberg 的解释成立；若不下滑，则本文成立。这一分界也为第七节（三）的 11 件判决提供了劳动力市场一侧的对应：法院的“由谁建立、记于谁账”与分析师市场的“由谁发起、向谁付费”，是同一条规则在两个制度里的两种写法。Burt（1992）的结构洞理论将网络位置本身视为资本；本文补充：它是谁的资本，取决于记录于谁的账上。《合同法重述（第二版）》第 188 条评注 g 所确立的“商誉归雇主”原则，本文借用其判法，但指出其只适用于名单由雇主记账的情形。

零知识证明。Goldwasser、Micali 与 Rackoff（1985）的零知识证明恰为“证明持有而不泄露”的方法，是本文在方法上最近的邻居。差别在于零知识证明要求秘密在两次证明之间保持不变，而前账阶段每次不同。对重复项而言，零知识证明是优于本文路径的答案。

一般技能例外的规范评价。Hrdy（2019）认为法院漏判了“技术上是秘密却不受保护”的信息类别，属于应予纠正的失误。本文借用其发现而结论相反：这不是失误，而是账本的结构——该类要素一旦入账即须归类且不可回复，法院拒绝将其记为秘密是正确的，只是理由被写成了失误。可判定的差别是：若依 Hrdy 的建议改进判定标准后该类信息能被稳定保护，其判断成立；若改进后仍被判为一般技能，本文成立。

上述文献共享一个默认：技能先于其入账而存在，入账不改变技能。Arrow 的交易对象、Teece 的资产、Becker 的资本、零知识证明的秘密、Zander 与 Kogut 的可编码性，均是先在彼处、而后被记录之物。一旦承认入账触发归类，这些概念便失去稳定的指称——

Arrow 的信息在被看见之前是什么, Teece 的资产在被记为资产之前是什么。本文的命题落在这一共同地基之上, 因而未被上述任何一支文献辨识。

(二) 适用边界

本文仅适用于仍有人为"下次在场"付费的技能。对于无在场者的技能——一次性产品、可完全由产物结算的工作——第三阶段归零, 其正确答案是零知识证明或商业秘密法而非本文路径。

本文的路径不得反向理解为"越不重复越保值"。技能习得研究给出了硬约束: 有功能的变异仅存在于任务流形之内 (Latash et al., 2002), 为变异而变异的解不是技能而是失误。前账阶段必须落在仍有目标的任务之内。

本文自身亦受命题约束。本文将三种传统的顶撞与判例中的发现写成可复述的三阶段路径并公开发表, 即完成了一次入账; 此后依此路径行事者所得为其快照。本文无法依靠被引用而保值, 这是命题 1 对本文自身的直接后果。

(三) 政策反作用

命题若被制度采纳, 最省事的实现方式有二, 且均会摧毁其所欲保护之物。其一, 设立"延迟入账"考核, 要求员工证明其保留了未入账部分——此举迫使"不入账"被写入账本。其二, 雇主在读到名单载体规则后, 将员工的每次客户接触强制录入客户系统——名单载体整体移入雇主账本, 员工自建名单的证据自此不存在。本文未能为二者找到制度性出口, 仅能指出: 任何以"入账次序"为依据对个人作出不利安排的机构(例如以"未留下文档"为由否定其贡献), 应公开编码规则并提供复核渠道; 安全关键系统(外科、航空、核电)的操作规程必须入账, 前账阶段仅存在于规程之外, 不得以本文为由拒绝书面化。

(四) 可检验预测

第一, 至 2028 年 12 月 31 日, 美国《保护商业秘密法》案件中, 被告仅凭记忆离职、未携带任何文件者, 法院判定争议内容为"一般知识、技能与经验"的比例不低于 70%。判定以公开判决书为准, 编码规则同附表; 机构或当事人自述不计; 样本少于 30 件则不裁决。若低于 70%, 本文第六、七节的经验推论为伪。

第二, 为名单载体另立预注册: 取以客户关系或客户名单为争点的公开判决, 样本不少于 40 件; 编码一项——名单由员工自建或早于雇佣, 抑或由雇主分配、记录或购入; 判负条款为两组判归差小于 30 个百分点; 判负后第四节(二)降为一项观察。预测: 至 2028 年 12 月 31 日, 名单由员工自建或早于雇佣的案件, 法院判归员工的比例不低于 80%。

本文的主张可分层引用: 第一层为"前账阶段"与"名单载体"两个概念; 第二层为表 1 的归属矩阵; 第三层为"法院以入账形式判归类"与"法院以记账方判名单归属"两项经验主张。第一层被推翻时后两层仍可独立成立, 最低成本的检验在第三层。

本文未能解释的一个现象是: 某些入账反而提升价值——论文署名、专利发明人、作品的作者名。这些入账没有将技能转入可复制区, 反而将其锁定于个人名下。一个可能的解释是它们记录的是"谁做的"而非"如何做的", 但本文没有材料证明这一点。

九、结论

本文从商业秘密法、表演研究与技能习得研究三个互不相容的解释出发, 指出其共享的前提——一切价值要素均可被无损记入账本——被商业秘密法自身百余年的判例所否定: 员工的一般知识、技能与经验一经写入雇主的账本, 即触发一次不可回复的归类判定。据此提出"前账阶段"概念, 区分账本的文件载体与名单载体, 并论证人工智能仅能吞并前者。33 件

判决的预注册检验支持归类判定随入账形式而异；11 件关系案件的探索性分析显示名单的归属随记账方而定，无一例外。

对个人的含义是明确的：在编码成本趋近于零的时代，技能保值不在于防止被写下——一切皆可被写下——而在于先于他人入账、仅入重复项，并使在场者名单记录于自己的账上。对法律的含义是：一般技能例外不是有待纠正的判定失误，而是账本结构的正确反映；商誉归雇主的原则应限于名单由雇主记账的情形。对人工智能治理的含义是：模型的能力边界与账本载体的边界重合，关于技能归属的争议将从“能否编码”转向“谁在记账”。

参考文献

Arrow, K. J. (1962). Economic welfare and the allocation of resources for invention. In **The Rate and Direction of Inventive Activity** (pp. 609-626). Princeton University Press.

Auslander, P. (1999). **Liveness: Performance in a Mediatized Culture**. Routledge.

Autor, D. H. (2014). Polanyi's paradox and the shape of employment growth. NBER Working Paper No. 20485.

Becker, G. S. (1964). **Human Capital**. Columbia University Press.

Bernstein, N. A. (1967). **The Co-ordination and Regulation of Movements**. Pergamon.

Bone, R. G. (1998). A new look at trade secret law: Doctrine in search of justification. **California Law Review**, 86(2), 241-313.

Burt, R. S. (1992). **Structural Holes**. Harvard University Press.

Cowan, R., & Foray, D. (1997). The economics of codification and the diffusion of knowledge. **Industrial and Corporate Change**, 6(3), 595-622.

Davids, K., Button, C., & Bennett, S. (2008). **Dynamics of Skill Acquisition: A Constraints-Led Approach**. Human Kinetics.

Deming, D. J. (2017). The growing importance of social skills in the labor market. **Quarterly Journal of Economics**, 132(4), 1593-1640.

Epoch AI (2024). **Open vs. closed AI: How behind are open models?** <https://epoch.ai/publications/open-models-report>

Fischer-Lichte, E. (2008). **The Transformative Power of Performance**. Routledge.

Goldwasser, S., Micali, S., & Rackoff, C. (1985). The knowledge complexity of interactive proof-systems. **Proceedings of the 17th ACM STOC**, 291-304.

Groysberg, B. (2010). **Chasing Stars: The Myth of Talent and the Portability of Performance**. Princeton University Press.

Hrdy, C. A. (2019). The general knowledge, skill, and experience paradox. **Boston College Law Review**, 60(8), 2409-2482.

Latash, M. L., Scholz, J. P., & Schöner, G. (2002). Motor control strategies revealed in the structure of motor variability. **Exercise and Sport Sciences Reviews**, 30(1), 26-31.

Lemley, M. A. (2008). The surprising virtues of treating trade secrets as IP rights. **Stanford Law Review**, 61(2), 311-353.

Meyer-Dinkgräfe, D. (2015). Liveness: Phelan, Auslander, and after. **Journal of Dramatic Theory and Criticism**, 29(2), 69-79.

Newell, K. M. (1986). Constraints on the development of coordination. In M. G. Wade & H. T. A. Whiting (Eds.), **Motor Development in Children** (pp. 341-360). Nijhoff.

Nonaka, I. (1994). A dynamic theory of organizational knowledge creation. **Organization Science**, 5(1), 14-37.

Phelan, P. (1993). **Unmarked: The Politics of Performance**. Routledge.

Polanyi, M. (1966). **The Tacit Dimension**. Doubleday.

Saunders, K. M., & Golden, N. (2018). Skill or secret? The line between trade secrets and employee general skills and knowledge. **NYU Journal of Law & Business**, 15(1), 61-120.

Schmidt, R. A. (1975). A schema theory of discrete motor skill learning. **Psychological Review**, 82(4), 225-260.

Shea, J. B., & Morgan, R. L. (1979). Contextual interference effects on the acquisition, retention, and transfer of a motor skill. **Journal of Experimental Psychology: Human Learning and Memory**, 5(2), 179-187.

Teece, D. J. (1986). Profiting from technological innovation. **Research Policy**, 15(6), 285-305.

Winter, S. G. (1987). Knowledge and competence as strategic assets. In D. J. Teece (Ed.), **The Competitive Challenge** (pp. 159-184). Ballinger.

Zander, U., & Kogut, B. (1995). Knowledge and the speed of the transfer and imitation of organizational capabilities: An empirical test. **Organization Science**, 6(1), 76-92.

判例 : Anderson Chemical Co. v. Green, 66 S.W.3d 434 (Tex. App.—Amarillo 2001); Cabela's LLC v. Wellman (Del. Ch. Oct. 26, 2018); Caudill Seed & Warehouse Co. v. Jarrow Formulas, Inc. (6th Cir. 2022); Combs v. Brick Acquisition Co. (Tenn. Ct. App. 2013); Corbin v. Tom Lange Co. (Tenn. Ct. App. 2003); CVD, Inc. v. Raytheon Co., 769 F.2d 842 (1st Cir. 1985); Dill v. Continental Car Club, Inc. (Tenn. Ct. App. 2013); Dur-A-Flex, Inc. v. Dy (Conn. 2024); Hanger Prosthetics & Orthotics East, Inc. v. Kitchens (Tenn. Ct. App. 2008); HCTec Partners, LLC v. Crawford (Tenn. Ct. App. 2022); Kewanee Oil Co. v. Bicron Corp., 416 U.S. 470 (1974); Oak Mortgage Group, Inc. v. Ameripro Funding, Inc. (Tex.

App.—Austin 2015); Total Staffing Solutions, Inc. v. Staffing, Inc., 2023 IL App (1st) 220533; Vantage Technology, LLC v. Cross, 17 S.W.3d 637 (Tenn. Ct. App. 1999); Variable Annuity Life Insurance Co. v. Miller (Tex. App.—Amarillo 2001); WorldVue Connect v. Szuch (5th Cir. 2025).

附表：33 件判决的逐件编码见随文文件。

作者贡献与声明：研究问题、理论框架与判定由作者提出；判例检索、编码与文稿整理使用了大型语言模型辅助，编码为单人完成，未经双人复核，作者对全部内容负责。作者声明无利益冲突。

十章读毕，剩下三件收尾的活，三件都是账。第十五章交**分界的账**：四十五对两两判定形式全表——其中十八对是十章正文里原生的，二十七对是本书补写的，逐条标〔补〕，标记本身就是第三章审计乙判负的兑付。第十六章交**检验的账**：十一场未跑的实验——十章各一场旗舰检验，加上裁决本书自身的第十一场异基底重跑；每场附判负条款与最低成本执行者。第十七章交**代价的账**：十章反噬预言的合流、三条伦理底线、与全书失效条件的索引。

合卷不做总结。该总结的第二章已经总结过；这三章存在的理由恰恰相反——把这本书交到能推翻它的人手里。

第七部 合卷：矩阵、纲领与失效

第十五章 四十五对分界

十个读数若有两个其实是同一个量的两个名字，收敛就掺了水。本章把四十五对全部摆平：每对给一条**判定形式**——一句可执行的观察或实验，说明在何种结果下两章分开、在何种结果下一章吸收另一章。

来源标注遵守第三章审计乙的裁定：标【原】的十八对，判定形式取自或压缩自十章正文（分界是各章自己长出来的）；标【补】的二十七对由本书补写——补写者与十章同产线，这些判定形式的独立性弱于原生分界，使用者应把它们当作待检的设计，而非已立的墙。

#	对	判定形式（一句）	源
1	现配优势 ↔零判份	同人同配组五年：显露表现稳定而判别力单调下降→零判份对；两者同步→现配吸收	〔原〕
2	现配优势 ↔脱稿步	现配锚人与场（交接后复现份额），脱稿锚产物与路（追问续签）；交接清单复现高而追问续签低的人存在→两轴独立	〔原〕
3	现配优势 ↔常席差	现配管载体（中断六月再复合，衰减重建则对），常席管归属（AI常驻后人效应方差压缩则对）；两预测可各自独立判负	〔补〕
4	现配优势 ↔单遍段	谈判既是当场配组又是单遍段：可传性轴与可重复性轴正交——同一任务可高s低Δ，亦可反之	〔补〕
5	现配优势 ↔峰兑	现配读每次交手的显露，峰兑读稀发事件的结算；平时交手读数与峰时兑付额相关低于阈值→两量独立	〔原〕
6	现配优势 ↔盲遍	配组当场做出（空间轴）与答案到达位次（时间轴）正交：清单复现率不随答案先后变，门位次不随配组轮换变	〔原〕
7	现配优势 ↔底流	配组不变、只改道低价值流量：现配预测显露稳定，底流预测无辅助底座衰退——一组人一年可裁	〔原〕
8	现配优势 ↔他手无 事账	固定配组只改取样预告期：录用排序翻动→他手对（时刻轴独立于配组轴）；不翻→他择率无增量	〔补〕
9	现配优势 ↔入账之 前	"写进清单即被批量供给"与"文件载体入账即入可复制区"同现象两读；判：对留种率低的优势执行"先自入重复项"处方，s是否保全——保全则两章分工成立	〔补〕
10	零判份↔ 脱稿步	零判份读产出与世界样本之差，脱稿读一次被发起的追问；档案同、未被追问者p有定义而q无定义→定义域即分界	〔原〕
11	零判份↔ 常席差	同批毕业生上未见率与归读率的相关<0.6→正交；≥0.6→两读数须合并重定义	〔补〕
12	零判份↔ 单遍段	零判份管产出的判别消耗，单遍管尝试的对象改写；可无限重投的任务上ρ照降而Δ≥0→两机制独立	〔补〕
13	零判份↔ 峰兑	未见率量"世界拿不拿得出"，盲读率量"量具分不分得开"；世界拿不出而量具照样分不开（判等差忽略吞掉真差）→两章各管一段	〔原〕
14	零判份↔ 盲遍	未见份高而门位次为零（原创产出但答案先看）可同时成立→暴露轴与位次轴正交	〔原〕
15	零判份↔ 底流	停喂者（β降）后续产出的p若加速下降→一因两果；β高而ρ照降（世界样本自行膨胀）→两机制可分	〔补〕
16			〔补〕

	零判份↔ 他手无事 账	同一作品集 p 高而他择率为零可并存；他择读数不改 p 、只改可信度→产出轴与 取证轴正交	
17	零判份↔ 入账之前	保密而已被本人私账记录之技能：零判份预测 p 照降（暴露即耗），入账预测归 类未触发（无第三方载体）——保密实验一次裁两章	〔补〕
18	脱稿步↔ 常席差	残差要人跨回合反复出现，续签要追问当场发生；档案厚而从未被追问者 g_r 可 读而 q 无定义→定义域不同；追问记录入字段后 g_r 是否变→耦合实验	〔补〕
19	脱稿步↔ 单遍段	单遍管对象不承认第二次，脱稿不缺记录、缺"持有者时间"字段；可重走任务上 的追问照样有效→两章不同轴	〔补〕
20	脱稿步↔ 峰兑	口试回溯两读：脱稿读"验证续签"，峰兑读"量具更替"；追问若可预告仍有效→ 脱稿侧；只有无预告才有效→峰兑让位他手	〔原〕
21	脱稿步↔ 盲遍	判决对照已写死：档案与追问全同的两人，形成期门位次分布不同则后续盲承诺 表现不同→盲遍对，脱稿对此无预测	〔原〕
22	脱稿步↔ 底流	每次成功续签同时是一笔三成分执行（真实后果、无辅助、亲手）；只靠追问喂 养能否维持底座→能，则底流"日常联合副产"降为充分非必要	〔原〕
23	脱稿步↔ 他手无事 账	预告的题库口试：脱稿仍认（锚定本人产物即可），他手降档为他择·定时→两 章对"预告"的敏感度不同，一场预告实验可分	〔补〕
24	脱稿步↔ 入账之前	两章同管记录缺口：脱稿要加"持有者时间"字段，入账警告凡入字段者皆文件载 体可被吞；若加字段后判别力恢复且不被模型学走→脱稿侧；恢复后随即失效→ 入账侧上限成立	〔补〕
25	常席差↔ 单遍段	常席管读数归属，单遍管重走结构；跨配组残差稳定者其单遍履历可厚可薄→两 量互不推出	〔补〕
26	常席差↔ 峰兑	常席要求配组轮换（空间条件），峰兑要求事件到场（时间条件）；轮换充分而 无事件者 g_r 可读 m 空账→定义域不同	〔原〕
27	常席差↔ 盲遍	跨配组残差与答案位次正交： g_r 高而 k 恒零（残差稳定但事事先看答案）者迁 移新题的表现=两章分水岭	〔原〕
28	常席差↔ 底流	互相误判点已写死：底流断绝者跨配组残差趋零，会被常席读作"本来就不存 在"——两读数须并置，残差消失的时间序（先停喂后趋零）归底流	〔原〕
29	常席差↔ 他手无事 账	配组轮换但全部预约，与配组固定但无预告抽查：两种记录各自独立改变 g_r 与 他择率→两轴正交	〔补〕
30	常席差↔ 入账之前	常席说能力读数不可裁归个人，入账的判例显示关系名单可按记账方裁归——能 力账与关系账是两本账；把客户名单当能力读数处即误读区	〔补〕
31	单遍段↔ 峰兑	全动模拟机是分方石：程序层可核校，峰兑认半个（承受层收窄），单遍不认 （对象未被耗散）——同一设备两章给出不同折扣	〔补〕
32	单遍段↔ 盲遍	答案先到而对象未动（纯判断）只触发盲遍；对象已改而无外部答案（纯操作） 只触发单遍——两门各自开合，四格可拆	〔补〕
33	单遍段↔ 底流	首过不可赎回（单遍）与底流可重建（停喂后复喂）：中断者复训一年，若底座 回升而首过缺口仍在→两章各管一段	〔原〕
34	单遍段↔ 他手无事 账	可重走任务上照样能做无预告抽样（ $\Delta \geq 0$ 亦可他择）→时刻轴独立于重走轴	〔补〕
35	单遍段↔ 入账之前	两个单向门：单遍在对象侧（私下重试也回不去），入账在记录侧（未入账即可 私下重试）；取一件可私下重试之技能即分	〔补〕

36	峰兑↔盲遍	峰时事件天然答案后到（现场无解析）， $k \geq 1$ 自动成立；分界在平时——盲遍对平时位次有预测（预走者迁移变差），峰兑对平时不作预测	〔补〕
37	峰兑↔底流	十章间最便宜的判决实验：同一名多年无事件在职者，峰兑判承接力仍在只是读不出，底流判底座已随停喂衰退——一次真实峰时事件加两组人同裁两章	〔原〕
38	峰兑↔他手无事账	峰时天然他择，但他择不必峰时（高频无事）；建立高频无预告小抽样账后，其存量若预测峰时表现→他手吸收峰兑一半；只有真实事件才预测→峰兑独立	〔补〕
39	峰兑↔入账之前	峰兑的兑付记录须对手方签收=一次由他方入账；入账警告签收件即文件载体可被复制——处方合成一条：兑付记录须名单化（记“为谁扛过”），两章各管一半	〔补〕
40	盲遍↔底流	β 高而 k 恒零者（天天亲手做但事事先看答案）迁移新题的表现=两章分水岭；底流已自认此处最可能被修正（“无辅助”或须细化为答案位次）	〔原〕
41	盲遍↔他手无事账	答案后到但取样自择，与答案先到但取样他择：四格可拆——位次读被读者与答案的关系，他择读被读者与读者的关系	〔补〕
42	盲遍↔入账之前	两个无痕的单向门：预走改认知态不留档，入账改记录态必留档；答案先到但未入账的一件交付——盲遍已触发、归类未触发	〔补〕
43	底流↔他手无事账	亲手（供给之手=被读者）与他手（取证之手=读者）是同一事件的两侧： β 高而他择率零（天天做无人看）与 β 零而他择率高（常被查全靠辅助）都存在→两轴独立且一事件可同入两账	〔补〕
44	底流↔入账之前	入账教“先自入重复项”，底流警告“代产即断供”；本人亲手产出后自入其重复项—— β 不降且前账保全→两处方兼容；他人代产代入→两章同判死	〔补〕
45	他手无事账↔入账之前	他手的第三条条件“归己”恰是入账的“首记于持有者账本”：他择读数若首记于机构账（如民航飞行数据）→空格不满且归类不由本人——两章在归己一格咬合而轴不同（时刻 vs 载体）	〔补〕

三点存目，纪律所需：

其一，第 8、16、23、29、34、38、41、43、45 九条涉《他手无事账》的补写分界之外，该章尚欠一段对**医学教育工作场评估**一族（受训者就真实临床任务被直接观察、结果记入本人档案）的点名：查其规程，受训者自选时机与观察者——落在空格上方一行（自择·他手·归己），不填充格，反而坐实“时刻归属从未被当作对象”；这段分界本应由该章自写，此处代记。**其二**，第 9、17、24、30、35、39、42、44、45 九条涉《入账之前的技能》的补写分界之外，该章尚欠同法域两部先行研究的谱系交代：员工知识归属的百年法律史（其核心问题——在公司可拥有的知识与自由人自留的知识之间划线——正是该章材料的主场）与径直以“人力资本法”命名此领域的一脉；该章的残余（入账触发归类、载体二分、记账方规则）不被二者覆盖，但谱系欠着就是欠着。**其三**，与本书十章同期尚有《能力证据债》《收方工序》《终稿之前》《空签》《弃项》等相邻研究，不收入本书，矩阵不含它们——它们与十章的分界，留给收它们的那本书。

第十六章 经验纲领：十一场未跑的实验

这一章是全书的欠条汇总。十章各自的证伪条款散在各章正文里；此处只收每章**最重的那一场**——判负即动摇该章承重命题的旗舰检验——外加裁决本书自身的第十一场。每场三栏：实验、判负条款、最低成本执行者。全部未跑；协议细目以各章正文为准。

场	章	实验（一句）	判负条款（一句）	最低成本执行者
1	盲遍	同质多案例随机化答案位次（先承诺后给答案 vs 先给答案），以迁移新题表现为结局	两组迁移无差且跨域稳定→位次层为伪	任何一门课的任课教师
2	单遍段	跨域重做样本库：同人同任务真实二次尝试的 Δ 审计	Δ 在申报的单遍域内系统性≥0→单遍性为伪	任一保存重考／重做记录的机构
3	现配优势	交接复现实验：完备清单交同资历者执行，N≥6，控文档质量	复现份额稳定≥九成并维持一年→当场性为伪	任一双人以上同岗团队
4	脱稿步	判者一致性预实验（约一百二十次追问，加权一致性功效计算已备）+既有面试录音回溯编码"叙述题 vs 延长线题"	一致性置信下界过低→清点规程整体重写	任一保存结构化面试录音的雇主
5	底流	改道实验：配组不变、仅把低价值流量改道给辅助系统，随访无辅助底座	底座不随停喂时距衰退→供给条款为伪	任一愿意分组的团队（与第 6 场合并执行）
6	峰兑	一次真实峰时事件、两组人（按事前兑付记录存量分组），事件窗口十二月自然等待，不得制造事件	事前存量不预测事件表现→峰兑为伪	任一有事件记录的组织——与第 5 场同一实验，一次裁两章
7	常席差	用双向固定效应公开复现数据，实测 AI 普及前后"人效应方差"与"跨组织人效应相关"的变动	方差未压缩且相关未升→常席差对读数分解传统的判决预测为伪	任一计量经济学研究组（数据公开）
8	零判份	保密条件检验：同质产出封存与公开两组，追踪 ρ 的收窄速度	保密组 ρ 不降→"暴露驱动"为伪，退回"使用驱动"	任一持有未发布产出库的机构
9	他手无事账	剂量-反应机制检验：控制成绩与作品集评分，入职前"他择·无预告读数份数"对入职十二个月绩效	无预告档系数不显著或为负→本章命题为伪	一家愿意回看自己招聘档案的机构（N≥30）
10	入账之前	名单载体四十件判例的正式预注册检验（该章已自立预注册）	记账方两组判归差小于三十个百分点→名单规则降为一项观察	任一法学院实证课程
11	本书	异基底重跑 ：同一道题、同一套方法，交与本产线无关的基底（另一模型或人类研究者）独立走完；三名盲评按事先注册的通约性规程判定其母命题与"发生差"的关系	判定为不可通约→本书收敛主张降级为产线伪影候选，改判写入再版首页	任何持有另一基底的团队

除这十一场之外，第四章另立**八条对外判决形式**——每条裁定的是本书与当代某一族说法的分界（清单派的效度背离检验、专长扩展派的三年随访、判断规则外化实验、增强对自

动化的两类岗位对比、参差前沿平滑后的随访、停喂时距对使用强度的独立性检验、代理工程作品集对他择读数的预测力对照，以及第十二族的品味半衰期测量）。它们与本表十一场的分工是：本表裁的是**本书对不对**，第四章裁的是**本书是不是别人已经说过的**。两类都未跑；第四章八条的执行人多半不是本书读者，而是那十二族自己——这一点无法强求，只能写在这里等。

三条纲领性说明。**其一**，第 5 与第 6 场是同一场实验的两个读数端，这是十章之间最便宜的一次裁决——设计一次、宣判两章；执行人优先做它。**这一场的完整预注册协议已写出，收在附录四**：含场景与备选、三角色隔离、两侧事前变量的编码规则、结局的去标识与一致性门槛、六条判负条款、灵敏度内对照与停止规则，另附一个两周可跑完的盲取回溯版。协议未执行；它进书不减少本表上的任何一张欠条。**其二**，十一场里有十场不需要本书作者参与，第十一场则**必须**排除本书作者参与——这不是姿态，是第三章三之四那处坦白的直接推论。**其三**，纲领的动机不是完备，是还账：一条论证“发生差要靠被检验的执行来喂养”的产线，自己的每一章都该有一次被真实对手方检验的机会。

第十七章 反噬、伦理与失效

十七之一 反噬的合流

十章各自的反噬预言，形状惊人地一致，值得合流写一次：**任何读数一旦成为考核目标，它所读的那件事就开始被制造**。留种率进考核，就有人表演配组、做低清单件成绩；盲态申报不可审计，就有人表演“不查答案”；峰兑进薪酬，就有人制造事件——纵火的消防员是现成的先例；追问一旦题库化，脱稿步自毁；无事一旦可累积成资产，就出现制造无事的市场——消费信贷已经演过一遍；无预告一旦可申报，就出现内部提前通气的“假他择”。这不是十个巧合，是同一条定律在十处的十次具体化：读数与被读物之间的因果，会被激励反接。

十章各自找到了半个出口——结算权外置、检查库自愿制、编码规则公开、把“不入账”排除出可考核项——合起来仍不是一个整出口。本书不假装有：**合卷也看不到完整的出口**。能写下的只有一条元规则：凡采用本书任何读数做分配的机构，必须同时采用该读数所属章的反噬条款，二者不得拆开引用。

十七之二 三条伦理底线

十章的伦理条款合并同类项后剩三条，全书通用，不得删减：

第一条，读数指证据的位置，不指人的价值。他择率为零的人只是没被在他择时刻读过； β 为零的人可能只是被剥夺了亲手的机会；跨配组残差不可读的人可能只是没得到过第二个班子。任何机构以本书读数作不利安排之前，必须先答“此人有没有被读的机会”——答不出，安排无效。

第二条，进入须自愿，且有对价。无预告取样、事件账、行踪式安排，只对自愿进入“检查库”者成立——报名带抽查条款的项目、接受带事件账的岗位。无对价、全覆盖的被观察义务，不是本书的处方，是本书点名反对的东西。安全关键系统（手术室、驾驶舱、电网）的既有规程高于本书一切建议：不得以“制造他择读数”为名注入任何真实风险。

第三条，编码规则必须公开，复核通道必须存在。“取样时刻由谁定”“这一步算不算无辅助”“这份记录归不归己”——每一个判定都必须有写在明处的规则与可走的申诉路。一个不公开编码规则的发生差账本，与它要替换的黑箱评分，是同样东西。

十七之三 失效条件索引

全书的可推翻处集中列示，免得散落各章成为免罪符：本书自身三条失效条件见第三章三之五（异基底条款、通道条款、可读条件必要性条款）；对外的三条对赌见第四章四之九（吸收之赌、定义域之赌、复检之赌），其中吸收之赌与定义域之赌是**本书被别人的工作作废**的两条路，复检之赌则规定本书不得为已被抹平的分界辩护；十章旗舰检验的判负条款见第十六章合表；各章正文另有细目证伪条款共计逾五十条，以各章为准。截至付梓：**已触发并照登的判负**——审计乙（本书，第三章）、常席差 P1、单遍段任务级检验、底流“无流量字段”强主张、脱稿步“标准无处可记”强主张与“追问不进字段”事实判断、入账之前“入账即归属”原表述；**自愿降格**——单遍段能力暴露侧；**未获裁决**——零判份两项、底流两次预注册。这张伤痕清单不是本书的弱点陈列，是本书唯一有资格自称的东西：每一处都是判负条款先于结果写死、结果不利照登的产物。

结语

回到那家餐馆。十把刀问完，好厨师的轮廓其实已经很老实：他在没看过菜谱的时候配出过菜；他手上有几个塌过就回不去的环节，而他都是第一遍就过的；他的好散在一手配组里，菜谱抄不走；拿着他的招牌菜让他当场续做，他续得上；他还天天亲手颠勺；换了班子他还是稳定好那一份；出事那天他在场，而且扛住了；密探不打招呼来过，没挑出毛病，那条记录他自己揣着；熟客名单记在他自己的本子上。——这十句话没有一句需要“人工智能”四个字才能成立，这正是要点：**AI 没有发明这道题，AI 只是把别的答案全部买断了。**

草坪上的那个缺口，本书用三章框架、两场审计、四十五条分界和十一张欠条，尽力证明它是踩出来的而不是画出来的。证明到什么程度，第三章已经如实标价：源独立，判据部分独立，生成不独立——最后一层的账，只有第十六章第十一场能清。

全书最诚实的一句话不是本书写的，是第九章写的，抄在这里作结：“**本产线问题十跑，无一次被真实有对手方的检验裁决过——这条产线产出量很高，而它的 β 很低。**”这本书没有资格例外。它交出去的这一刻，是一份自择证据：写作时刻是作者选的，交付时是准备好的。它能获得的第一笔他择读数，是某位读者在作者不选择的时刻、按第十六章的某一场协议，真的去跑——判负也算入账，而且尤其算。

给读者留一句可以随身带走的判据，不含任何“应当”：

下一件你交出去的东西——答案是什么时候到的？时刻是谁选的？记在谁的账上？

三个问句答完，你名下的发生差，比任何清单都诚实。

参考文献

体例说明。本书第五至十四章为十项独立研究的原文，各章末尾保留其自身的参考文献与材料来源，按原样收入、未作合并——合并会掩盖一件本书需要读者看见的事：十项研究的文献池互不重叠，这正是第三章审计甲所测的源独立性在文献层面的投影。本页只列框架卷（第一至四章、第十五至十七章、结语与四附录）所引文献，按出现次序排列。网络来源的访问日期均为 2026 年 9 月 3 日。

一、收敛作为证据形态（第三章）

1. Whewell, W. *The Philosophy of the Inductive Sciences, Founded Upon Their History*. London: John W. Parker, 1840. (归纳的会通)
2. Levins, R. "The Strategy of Model Building in Population Biology." *American Scientist* 54(4), 1966: 421-431. (独立谎言之交集)
3. Wimsatt, W. C. *Re-Engineering Philosophy for Limited Beings: Piecewise Approximations to Reality*. Cambridge, MA: Harvard University Press, 2007. (稳健性分析)
4. Campbell, D. T., & Fiske, D. W. "Convergent and Discriminant Validation by the Multitrait-Multimethod Matrix." *Psychological Bulletin* 56(2), 1959: 81-105.

二、当代诸说（第四章，十二族）

1. World Economic Forum. *Future of Jobs Report 2025*. Geneva, January 2025.
2. 曾湘泉等：《人工智能时代，我们如何更好就业创业》，《人民日报》理论版，2026 年 4 月 9 日。
3. 《全球数字教育发展指数（2026）》，2026 世界数字教育大会发布；相关报道见《人民日报》2026 年 5 月 31 日第 5 版。
4. Autor, D. H. "Applying AI to Rebuild Middle Class Jobs." NBER Working Paper No. 32140, February 2024.
5. Agrawal, A., Gans, J., & Goldfarb, A. *Prediction Machines: The Simple Economics of Artificial Intelligence*. Boston: Harvard Business Review Press, 2018.
6. Agrawal, A., Gans, J., & Goldfarb, A. *Power and Prediction: The Disruptive Economics of Artificial Intelligence*. Boston: Harvard Business Review Press, 2022.
7. Agrawal, A., Gans, J., & Goldfarb, A. "Exploring the Impact of Artificial Intelligence: Prediction versus Judgment." *Information Economics and Policy* 47, 2019: 1-6.
8. Brynjolfsson, E. "The Turing Trap: The Promise & Peril of Human-Like Artificial Intelligence." *Dædalus* 151(2), Spring 2022: 272-287.

9. Dell'Acqua, F., McFowland III, E., Mollick, E., Lifshitz-Assaf, H., Kellogg, K., Rajendran, S., Kraye, L., Candelon, F., & Lakhani, K. R. "Navigating the Jagged Technological Frontier: Field Experimental Evidence of the Effects of Artificial Intelligence on Knowledge Worker Productivity and Quality." Harvard Business School Working Paper 24-013, September 2023 ; 正式发表于 *Organization Science*, March 2026.
10. Kosmyna, N., Hauptmann, E., Yuan, Y. T., Situ, J., Liao, X.-H., Beresnitzky, A. V., Braunstein, I., & Maes, P. "Your Brain on ChatGPT: Accumulation of Cognitive Debt when Using an AI Assistant for Essay Writing Task." arXiv:2506.08872, June 2025 (v2, December 2025).
11. Stanković, M., Hirche, E., Kollatzsch, S., & Doetsch, J. N. 对上文的方法学评论。arXiv, 2026.
12. Gerlich, M. "AI Tools in Society: Impacts on Cognitive Offloading and the Future of Critical Thinking." *Societies* 15(1), 2025.
13. Mollick, E. 关于"品味"的访谈言论, 2026年4月30日播出; 相关报道见2026年5月。
14. Karpathy, A. "Agentic Engineering." Sequoia Ascent 2026 演讲及作者讲稿摘要, 2026年4月30日。
15. 英国等地劳动力市场"向下拉平"机制的一族工作论文及其对 Autor (2024) 的评述。
16. Spence, M., "Job Market Signaling," *Quarterly Journal of Economics* 87(3), 1973: 355–374.
17. Collins, R., *The Credential Society*, Academic Press, 1979.
18. Messick, S., "Validity," in *Educational Measurement*, 3rd ed., Macmillan, 1989.
19. Kane, M. T., "Validating the Interpretations and Uses of Test Scores," *Journal of Educational Measurement* 50(1), 2013.
20. Ericsson, K. A., Krampe, R. T., & Tesch-Römer, C., *Psychological Review* 100(3), 1993.
21. Klein, G., *Sources of Power*, MIT Press, 1998.
22. Acemoglu, D., & Restrepo, P., *American Economic Review* 108(6), 2018 ; 及 *Journal of Economic Perspectives* 33(2), 2019.

三、实验方法学（附录四）

1. International Council for Harmonisation. *ICH E10: Choice of Control Group and Related Issues in Clinical Trials*. 2000. (灵敏度与内部对照)
2. Prasad, P., Sun, J., Danner, R. L., et al. "Excess Deaths Associated With Tigecycline After Approval Based on Noninferiority Trials." *Clinical Infectious Diseases* 54, 2012: 1699–1709. (本书第十一章据以推翻其共同前提, 此处为附录四灵敏度装置的来源)

四、本书自身的登记文件

1. 《发生差》框架卷两场审计·预注册，2026年9月3日。SHA-256：48a34a45723cf56923044599298ca284f817040ecdfe73c3c94e33c2e08904be。全文见附录一。
2. 《峰时对照：一次实验裁两章》预注册协议 v1.0，2026年9月3日。正文段SHA-256：
16e6c87277fd4de9deb144aaac8f91ad18288b126864e8c190f304e64b050339。全文见附录四；独立存档件为登记正本。

一处必须写明的缺项。第十五章其三与第四章四之九各自登记了本书尚欠的谱系交代：医学教育工作场评估一族之于第十三章，员工知识归属的百年法律史与"人力资本法"一脉之于第十四章。这两处的文献本页不列——**欠着就是欠着，列出来反而像已经交手过。**

附录一 两场审计预注册全文

以下为第三章两场审计的预注册原文，写于任何计数之前；其 SHA-256 摘要为 48a34a45723cf56923044599298ca284f817040ecdfe73c3c94e33c2e08904be，正文第三章所引与此一致。

《发生差》框架卷两场审计·预注册（写于任何计数之前，2026-09-03）

审计甲：源独立性

- 单位：十篇 $\times$ 三个领域 = 30 个源位。
- 编码规则（读取前写死）：两个源位记为「同一源」，当且仅当二者引用同一判据传统的奠基文献或同一监管文书体系（例：同引 UTSA/Kewanee $\rightarrow$ 同源；同属生态学但一引 BEF 超产、一引中性理论 $\rightarrow$ 不同源）。存疑一律偏严（记同源）。
- 判负条款：去重后唯一源数 $< 24/30$ (80%) $\rightarrow$ 「十路独立收敛」降级为「十路准独立收敛」，全书不得再以「独立」二字承重。
- 不利处置：复用清单照登全表，逐处注明复用的是同一判据之刀、还是同一田地的不同判据。
- 停止规则：30 位编码完即停，不因结果改编码规则。

审计乙：原生分界覆盖率

- 单位： $C(10,2) = 45$ 对。
- 判定「该对已有原生分界」当且仅当：至少一篇的正文（十篇原文，非本书新写部分）出现另一篇的篇名或读数名。别名对照表先写死：现配优势|留种率；零判份|未见率|未见份；脱稿步|脱稿成立率|可续签的署名；常席差|归读率；单遍段|重走落差|单遍力；峰兑|盲读率；盲遍|门位次|预走；底流|在喂率；他手无事账|他择率|他择存量；入账之前|前账阶段|名单载体。
- 判负条款：覆盖 $< 27/45$ (60%) $\rightarrow$ 「十篇在正文里两两互有分界」这一强主张为伪，降级为「部分原生可分」；缺失对由本书第十四章补写、逐条标〔补〕，并须在第三章明写：补写者与十篇同产线，其独立性弱于原生分界。
- 停止规则：45 对判完即停。

附录二 十读数速查手册

本表供执行者速查，每项五栏：读数定义、定义域（在哪读才有定义）、一条判据问句（可直接拿去问一份简历或一个岗位）、判负条款（怎样算这个读数错了）、反噬警戒（进考核后最先被制造的是什么）。细目以各章正文为准；本表与正文冲突处，以正文为准。

读数	定义（一句）	定义域	判据问句	判负条款（旗舰）	反噬警戒
门位次 k （盲遍）	答案到达落在本人成形承诺序列的第几步	答案可考时序的具体未知	这东西做成之前，他见没见过答案？	随机化位次实验迁移无差	表演"不查答案"、盲态申报造假
重走落差 Δ ／门位 k （单遍段）	重做样本上二次与首次成功率之差；工序中首个不可撤销步的位次	有真实重做记录的工序段	他的关键环节，哪一步走过就回不去？他第一遍过了吗？	申报单遍域内 Δ 系统性 ≥ 0	伪造"首过"、隐藏重投史
留种率 s （现配优势）	完备交接清单交同资历者照做后，接收方能复现的份额	可交接、有接收方的真实交付	他的好处，抄走清单能复现几成？	复现份额稳定 ≥ 9 成并维持一年	表演配组、故意做烂清单件
脱稿成立率 q （脱稿步）	在本人产物延长线上被无脚本追问时，续签成立的比率	一次真实被发起的追问	拿着他的作品当场往下问，他续得上吗？	判者一致性置信下界过低	追问题库化、押题式排练
在喂率 β （底流）	过去九十天内，能力清单条目发生过无辅助、有后果、亲手执行的比率	有真实赌注的日常产出流	他现在还亲手做吗——最近一次是什么时候？	底座不随停喂时距衰退	表演亲手、伪造无辅助日志
归读率 g_r （常席差）	人固定、配组轮换下残差符号一致的占比	跨配组的多回合记录	换了班子，还稳定偏向他的那份有多大？	人效应方差未随 AI 常驻压缩	挑配组、避强同场者
盲读率 m （峰兑）	峰时事件整笔结算额与平时量具读数之比的缺口	谁也没选的稀发事件处	出事那天他在场吗？扛住了吗？记在他名下了吗？	事前兑付存量不预测事件表现	制造事件、抢占事件署名
未见率 ρ （零判份）	产出中世界尚不能复现同类物的份额	相对世界现存样本、随暴露收窄	这份产出里，外面拿得出来的占几成？	保密组 ρ 不降（暴露驱动为伪）	囤积保密、拒绝发布以护 ρ
他择率 $\times$ 归己率（他手无事账）	证据中"时刻非本人所选、他人之手取得、无异常、记入本人可携账本"的份数	他择时刻的取样记录	他的证据里，有几份的时刻不是他自己定的？	无预告档对入职绩效系数不显著	假他择（内部通气）、造无事
记账方判据（入账之前）	技能最近一次入账由谁写入；名单载体首记于谁的账本	未入账段与名单载体	他的命根名单，记在谁的本子上？	记账方两组判归差 < 30 个百分点	抢先代入账、名单圈占

十栏"判据问句"连读，就是结语那三问的展开式。执行者若只带走一页，带走这一页；若只带走一句，带走那三问。

附录三 三类读者的用法

全书到此为止讲的是"差在哪里才有定义"。这一附录讲"那么我该做什么"。三类读者各一节，每节三样：**先改哪一处**（成本最低、见效最快的一刀）、**怎么记账**（要加的字段，不要加的字段）、**不得为**（与第十七章三条伦理底线绑定，不得拆开引用）。

三节共守一条前提，写在最前面：**本书的读数目前一场旗舰检验都没跑过**（第十六章）。因此这一附录里的任何一条，都应当以试点、可回撤、可申诉的方式采用，不得直接进考核、进薪酬、进录用红线。谁把它当成即插即用的评分表，谁就正好落进第十七章第一条反噬。

一、用人单位

先改哪一处：把一次自择取样换成他择取样。不是加流程，是**替换**——把原本"请候选人回去准备一份作品／方案"的那一环，换成"在双方都没准备的时刻，就他已经交来的东西当场往下问二十分钟"。成本几乎为零：不加入、不加轮次、不买系统，只把同一时间段从"准备好的东西"挪到"看没准备的时刻"。

第一个月就能看到的差别有两种。一种是候选人之间的排序翻动——翻动本身不证明新方法更好，但它证明旧读数没在读你以为它在读的东西（《他手无事账》）。另一种是有些人在追问的第三步之后开始复述而不是往前走——那正是《脱稿步》说的续签失败，档案上完全看不出来。

怎么记账。招聘与绩效档案里加四个字段，都极短：

字段	记什么	为什么
取样时刻由谁定	本人／对方／随机	排除排演，四格空位的第一格
答案到达位次	承诺在前／答案在前／不可考	这份产出是不是这个人自己的一遍（《盲遍》）
无辅助亲手执行	最近一次的日期	底座还在不在被喂（《底流》）
配组构成	同场者与工具	事后可做跨配组残差（《常席差》）

不要加的字段同样重要：**不要记"AI 使用时长／次数"**。它读的是使用强度，而本书的驱动变量是停喂时距（第四章四之七）；记错变量会把一个每天亲手做事、也频繁用工具的人误判为退化。

不得为。无预告取样只对**自愿进入检查库**者成立，且须有对价（岗位、项目、优先权），不得全员覆盖、不得溯及既往；编码规则（什么算无辅助、什么算他择）必须公开并留申诉通道；不得以制造他择读数为名在安全关键岗位注入任何真实风险。一个不公开规则的发生差账本，与它要替换的黑箱评分是同样东西。

二、教师与培养方

先改哪一处：把"交作业"改成"交作业＋一次不预告的当场续做"。同一份作业，同样的分值，只把评分的重心从成品挪到延长线上的三分钟。这一刀之所以最便宜，是因为它不需要判断学生用没用 AI——它绕开了那个已经无法裁定的问题：不问"这是不是你写的"，问"就着它往下走一步"。走得动，署名有效；走不动，成品再漂亮也只是收据。

第二刀成本略高但值得：**在一门课里留出一处"答案先不给"的工序。**哪怕只有一次作业、二十分钟，要求先交出成形的承诺（一个判断、一个方案、一个预测），再放开检索与工具。这一遍对每个学生每个具体未知只有一次（《盲遍》），错过就没了；而它是整门课里唯一能证明"这个人在没有答案时能走到哪"的记录。

怎么记账。学生的成长档案里，值一记的不是分数曲线，是三样：承诺在前的次数、无辅助亲手完成的最近日期、被无预告追问后续签成立的比率。三样都不需要新系统，一张表格即可。教师最该警惕的一件事写在《底流》里：**把低价值的重复练习整批交给工具那天，成绩、作业质量、教学进度三本账会同时改善，而学生的底座在无声地断供。三账全绿正是断供的样子，不是断供的反证。**

不得为。不得把这些读数用于分班、评优或任何不可回撤的分流；不得在未成年人身上做无预告取样而不告知其监护人；不得把"续签成立率"变成一个可以刷的指标——一旦追问题库化，《脱稿步》当场自毁（第十七章十七之一）。

三、毕业生本人

给你的不是励志话，是三件今天就能做、且成本你自己承担得起的事。

第一，给自己留一份他择证据。你名下所有能拿出手的东西，几乎都是你选好时刻交出去的：简历、作品集、项目复盘。它们在准备可以外包的时代同时贬值。去弄到哪怕一份不是你选时刻的记录——一次无预告的当场答辩、一次临时顶上的活、一次别人挑的时间做的评审。一份就够，因为现在几乎没有人有。

第二，别把亲手做的那部分整批交出去。这一条最反直觉，因为交出去时三本账全绿：你更快了、产出更好了、评价更高了。请记住《底流》的读数——过去九十天里，你的能力清单上有几条**发生过至少一次无辅助、有真实后果、亲手完成的执行**？只要这个数在掉，其余三本账再好看都不作数。也别把它做成苦行：不是拒绝工具，是每周留几件真有赌注的小事自己从头走完。

第三，看住那本记名单的账。你写下的一切"怎么做"，迟早进入可复制区；进不去的是"为谁做过、谁认得你、谁欠你一次"（《入账之前的技能》）。判据只有一句：**你的关系名单，首记在谁的本子上？**记在公司系统里的，公司拿走；记在你自己本子上的，才跟着你走。

最后一句，与第十七章第一条伦理底线互为表里，请连着读：**这三条能不能做，很大程度上取决于你有没有被读的机会**——有没有人愿意在不预告的时刻读你，有没有第二个班子让你出现，有没有一处让你亲手做事的位置。机会不平等的时候，读数上的差首先记录的是机会的分布。所以这三条对你只是建议，对任何机构都不是判人的许可：**读数指的是证据的位置，不是人的价值。**

附录四 峰时对照：一次实验裁两章

编者说明。本附录收录的是《底流》（第九章）与《峰兑》（第十一章）联合判决实验的预注册协议全文，即第十六章表中第 5 场与第 6 场合并后的那一场。收录它有三点须先讲明，免得读者把它读成本书的成果：

其一，它未执行。本书付梓时，这场实验一次也没有跑过。协议进书不减少第十六章那张欠条表上的任何一张欠条——十一场仍是十一场，只是其中一场从"一句话的设想"变成了"一份可存档、可复算、可被别人拿去执行的文件"。这两者之间的差别是实在的，但它不是结果。

其二，书内是复制件，登记以独立文件为准。该协议以独立文件存档并公开其 SHA-256 摘要；本附录为便于阅读做了标题层级的调整，正文一字未改。凡校验哈希者，请以独立文件为准；两处若有出入，以独立文件为准，本附录作废。

其三，本书作者不得执行它。协议第二节之三写死了角色隔离：撰写方不得担任事前变量编码员、结局评定员或分析员中的任何一角，不得接触任何一侧数据。这一条与第三章三之四那处坦白同源——一条自认"生成不独立"的产线，不能同时是自己的裁决者。因此这份协议放在这里，是一份**交出去的东西**：它的价值全部取决于有没有人愿意在本书作者不在场的情况下，真的去跑一次。

被检的两章 《底流》（本书第九章）与《峰兑》（本书第十一章）

协议撰写 王德生／德麦国际出版社 · SDE 学派，与大语言模型 Claude 协作

日期 2026 年 9 月 3 日

状态 未执行。本文件为协议本身，写于任何数据接触之前。

登记方式 本文件全文哈希公开（见第十二节），执行方在取数前存档，取数后不得修改。

一、这场实验要裁什么

《底流》与《峰兑》对**同一个人**给出方向相反的判断，两章都把这一点写在各自正文里。

取一名多年没有遇到过真实高压事件的在职者：

- **《峰兑》判：**他的承接力仍在，只是平时量具读不出——期内读数为零是这份价值的**构成性质**，不是它不存在的证据。
- **《底流》判：**他的底座已随停喂时距衰退——在这里，"读不出"与"不存在"正在合流。

两说都能解释"平时看不出差别"这件事，因此在平时不可分。它们只在一个地方分开：**当真实事件终于到来时，谁的表现更好，由什么预测。**

- 若表现由**事前的兑付记录存量**预测，而与停喂时距无关 → 《峰兑》对，《底流》那句"读不出即在漏掉"为伪。
- 若表现随**停喂时距**单调下降 → 《底流》对，《峰兑》第一层（期内盲是构成性质而非衰退）在本场景为伪。

一次真实事件、两组人、两个事前变量，同时裁两章。这是全书十一场检验里唯一一场一次实验裁两章的，也是唯一一场不需要新建任何数据采集系统的——它吃的是组织已经在写的事故档案与排班表。本书第十六章因此把它列为执行者的第一优先。

必须先写死的一件事：本协议**不制造事件**。事件必须是自然发生的、本可以不发生的、且与本研究无因果关系的。凡为获得读数而制造、放任或延迟处置风险者，本协议自动作废，且触及本书第十七章第二条伦理底线。

二、场景与执行者

2.1 主场景（推荐）：技术组织的重大故障响应

选它的理由有四条，都与“便宜”直接相关：**事件自然发生**（无需制造）、**时间戳精确**（告警、升级、恢复三个时刻都在系统里）、**留有署名档案**（事故复盘报告写明谁做了什么）、**存在真实对手方**（受影响的用户与客户，其利益不由响应者支配）。

- **事件定义：**组织内既有严重度分级中的最高两级（通常记作 Sev-1/Sev-2 或 P1/P2）故障，且满足：影响到组织之外的真实用户；发生时不存在回滚以外的“重来按钮”；有事复盘文档。
- **单位：**一次事件中的一名响应者，记为一个观测（responder-event）。
- **样本：**组织内所有进入该事件响应的人员，不得由任何人挑选。

2.2 备选场景与字段映射

场景	事件	"无重来按钮"的判据	表现档案
医院快速反应/急救团队	一次真实院内急救呼叫	患者状态不可回滚	抢救记录与复盘
客服/运营峰荷	一次真实峰荷日（大促、政策截止日）	当日窗口关闭即结束	工单与升级记录
电网/调度	一次真实故障或极端负荷	供电不可事后补发	调度日志
field service/检修	一次非计划停机抢修	停机损失实时发生	工单与验收记录

四个备选与主场景共用同一套变量与判负条款，只替换事件定义与档案来源。**不得使用演练、模拟、桌面推演或全动模拟机作为事件**——那些有重来按钮，按《峰兑》第四节第二条性质与《单遍段》的分界，不构成峰时结算。

2.3 执行者与角色隔离

最低配置为**三个人**，且三个人之间必须做角色隔离：

角色	做什么	不得接触
甲：事前变量编码员	事件发生 之前 ，清点每位在册人员的 β 、停喂时距 d 、兑付记录存量 S	任何事件表现数据
乙：结局评定员（ ≥ 2 人）	事件发生 之后 ，仅凭事故档案对响应者表现评分	甲的全部编码结果、当事人身份（尽可能匿名化）
丙：分析员	到达停止规则后合并甲乙两侧数据、按本协议跑分析	合并之前不得看任何一侧的分布

本协议的撰写方（王德生与本产线）不得担任甲、乙、丙中的任何一角，不得接触任何一侧数据。这是本书第三章三之四那处坦白的直接推论：一条自认"生成不独立"的产线，不能同时是自己的裁决者。协议作者可回答规程性询问，问答记录须与数据一并公开。

三、事前变量（甲方编码，事件之前完成）

3.1 《底流》侧：停喂时距 d 与在喂率 β

按《底流》正文的定义执行，不作改动。

三成分判据（缺一即不计为一次执行事件）：① **无辅助**——判据性步骤由本人走出，允许使用工具做检索与整理；② **真实后果**——搞砸了有人真的受损，损失不由本人支配；③ **亲手**——本人实际执行，不是审阅、批准或指派。

能力清单：对每位在册人员，由甲方按其岗位说明书列出 6-12 条能力条目，清单在事件发生前定稿并封存，此后不得增删。

- **停喂时距 d_i** ：第 i 条能力距最近一次三成分齐备执行的天数。
- **在喂率 β** ： $d_i \leq 90$ 天的条目占比（ $T = 90$ 天，沿用《底流》默认，该窗口来源已在该章正文交代）。
- **本场景主用量**：与本次事件直接相关的那几条能力上的 d_rel （取其中位数），以及全清单的 β 。

编码来源：排班表、提交记录、工单归属、值班日志——全部为组织既有记录，不得依赖自述。凡无法从记录判定的条目标记为"不可判定"，不得猜；不可判定条目超过清单一半者，该人整体剔除（剔除数须报告）。

3.2 《峰兑》侧：兑付记录存量 S

按《峰兑》第十节"可核验的兑付记录"四条构成执行，缺一条即不计入 S ：

① **真实对手方**（事件另一头站着利益不由本人支配的人）；② **真实赌注**（当时有不可撤销的损失在场）；③ **无重来按钮**（当时不存在补考、延期、回滚的选项）；④ **对手方签收**（事后留有可被第三人查证的结算痕迹）。

- **S** ：该人事件发生日之前、过去 36 个月内四条齐备的兑付记录条数。
- **S_recent** ：其中最近 12 个月内的条数（用于第七节的次级分析）。

《峰兑》正文已写明：这一栏在现行任何人事系统里都不存在，只能事后由结算处倒推。本协议因此规定： S 由甲方从事故复盘档案、客户往来、结算单据中 **逐条重建**，每条须附可查证据链；重建过程与判据须留痕并随数据公开。 **S 的重建必须在事件发生之前完成**——事后重建会被本次事件的表现污染。

3.3 共线性申报

d_rel 与 S 有可能高度相关（常出事的人也常亲手干活）。若二者的 **Spearman 相关绝对值 ≥ 0.5** ，本轮申报为"不可识别"，两章均不得据此判负，只报告相关本身。这条先于数据写死，是为了防止事后把共线性读成"两说其实是一回事"。

四、结局变量（乙方编码，事件之后完成）

主结局：事件内标准化处置表现 y 。

由至少两名评定员，仅凭事故档案（时间线、操作记录、复盘文档），对每位响应者在该次事件中的处置质量做 0-100 评分，评分维度事先写死三项：**判断的准确性**（关键判断是否正确且及时）、**动作的有效性**（所做操作是否推进了止损）、**升级与协作的恰当性**（该叫人时是否叫人、该停手时是否停手）。三项等权。

- **匿名化：**档案交付评定员前，由第三方去除姓名、工号、职级与可识别的个人称谓；无法完全去标识的事件须申报。
- **一致性门槛：**两名评定员评分的组内相关系数（ICC） ≥ 0.60 方可入库；低于 0.60 的事件整例剔除，剔除数须报告。若全库 ICC < 0.60 ，本轮判**未获裁决**（见第六节）。
- **事件内标准化：** y 在**每次事件内**做 z 变换。这一步是为了吸收事件难度的差异——不同事件的绝对难度不可比，同一事件内的相对表现可比。

次级结局：事件恢复时长中该人负责区段的耗时（若档案可拆分）；是否发生二次故障。次级结局不参与主判负。

五、假设与主分析

5.1 模型

混合效应线性模型，观测单位为 responder-event：

$$y_{ij} = \alpha_j + b_1 \cdot \log(1 + d_{rel,ij}) + b_2 \cdot S_{ij} + \gamma \cdot X_{ij} + u_i + \varepsilon_{ij}$$

其中 α_j 为 **事件固定效应**（吸收事件难度、时段、系统差异）， u_i 为 **人的随机截距**（同一人可能出现在多次事件中）， X 为事先写死的协变量：工龄（年）、职级（序数）、是否当班主责（0/1）、到场时距（分钟）。**协变量清单在取数前封存，不得增补。**

5.2 三条主假设

- **H1（《底流》侧）：** $b_1 < 0$ ——停喂时距越长，事件中表现越差。
- **H2（《峰兑》侧）：** $b_2 > 0$ ——事前兑付记录存量越大，事件中表现越好。
- **H3（分辨）：**在同时纳入 d_{rel} 与 S 的模型中，两者各自的偏效应量（偏 R^2 ）报告到小数点后三位，并给出二者之比。

显著性水平 $\alpha = 0.05$ ，双侧。**最小可探测效应（MDE）先行写死：** $|b_1|$ 对应"停喂时距每增加一个自然对数单位、表现下降 0.25 个事件内标准差"； b_2 对应"每多一条兑付记录、表现上升 0.20 个事件内标准差"。

5.3 灵敏度内对照（本协议的关键装置）

《峰兑》的核心论证是"量具会在灵敏度意义上死亡，而且死得无声"。这条论证反过来对本实验构成威胁：如果本实验的结局量具本身分辨不出人，H1 与 H2 都会读作"不显著"，而那不是两章都错，是量具死了。

因此本协议移植《峰兑》自己的娘家学科（非劣效试验方法学，ICH E10）的**灵敏度检验**：事先指定一组**已知应当有差的对照**——同一批事件中，**入职 < 6 个月者与同岗位工龄 ≥ 5 年者的 y 之差**。

灵敏度门槛：该已知差的估计值须 ≥ 0.35 个事件内标准差，且 95% 置信区间下界 > 0 。**未达门槛者，本轮整体判"无灵敏度 → 未获裁决"**，H1/H2 的结果一律不得解读，更不得作为任一章的判负依据。

这一条是本协议里最容易被执行者省掉、也最不能省的一条：没有它，一场无效的实验会被读成"两章都被证伪"。

六、判负条款（先于数据写死）

条款	触发条件	处置
甲：《底流》F1 判负	灵敏度门槛通过，且 b_1 的 95%CI 完全落在 $[-MDE, +\infty)$ （即效应不存在或方向相反）	《底流》"读不出即在漏掉"的判断在本场景为伪；该章第十五节其九的判决对照按不利结果公布，正文相应段落须改写而非删除
乙：《峰兑》第一层判负	灵敏度门槛通过， b_2 不显著为正，且 b_1 显著为负	《峰兑》"承接力仍在、只是量具读不出"在本场景为伪；期内盲不再可被主张为构成性质
丙：两说共负	灵敏度门槛通过， b_1 与 b_2 均不显著	两章在本场景同判"无预测力"，须由第三机制解释；本书第二章的存续条款与事件处条件各降一级
丁：各管一段	b_1 显著为负且 b_2 显著为正，且第三节共线性申报未触发	两章均获支持且互不吸收；须报告偏 R^2 之比，写入两章再版
戊：不可识别	d_rel 与 S 的相关系数绝对值 ≥ 0.5	本轮不裁决任何一章，只公布相关系数与散点
己：无灵敏度	灵敏度门槛未达	本轮判未获裁决；量具须重设后另立预注册

不利结果处置：无论触发哪一条，结果照原样公布，包括对两章都不利的丙条。执行方不得因结果不利而不发表；协议作者不得因结果不利而修改本协议或另立"更合适"的检验充数——《底流》正文已有两次"未获裁决"照登的先例，本协议沿用同一纪律。

七、样本量、功效与停止规则

目标样本： $N \geq 80$ 个 responder-event 观测，且来自 ≥ 15 次独立事件，且事件内至少 3 名响应者。

功效说明：在事件固定效应吸收事件间方差、组内相关约 0.2 的假设下， $N = 80$ 对 5.2 节写死的 MDE 具有约 0.80 的功效（双侧 $\alpha = 0.05$ ）。功效计算脚本随协议存档，参数不得事后调整。

停止规则（先到为准）：① 达到 $N = 80$ ；② 自协议存档之日起满 24 个月；③ 组织退出。到达停止规则时若 $N < 40$ ，判**未获裁决**，不得延长窗口"再等一等"——延长窗口等于按结果调整样本。

窥视纪律：达到停止规则之前，甲乙两侧数据不得合并，任何人不得计算 $H1/H2$ 的估计值。中期只允许查看**入库计数**（观测数、事件数、剔除数），不得查看任何分布或相关。违反窥视纪律须在发表时申报，且本轮降级为探索性。

八、最低成本版本：盲取回溯

上节的前瞻版需要等 12–24 个月。对已有 24 个月完整事故档案的组织，本协议提供一个可在**两周内跑完**的回溯版本，代价是证据强度降一级。

规程如下，次序不可颠倒：

1. **先封结局：**由第三方从档案库中抽出全部符合事件定义的事件，去标识后交乙方评分。乙方在完成全部评分并封存之前，不得接触任何人员的排班、提交或兑付记录。
2. **再编事前变量：**甲方在不知道任何评分结果的前提下，按第三节规则从档案中重建每位响应者在**该事件发生日之前的** d_{rel} 、 β 与 S 。所有编码只允许使用事件发生日之前的记录——甲方须逐条记录数据截止日，越界即剔除。
3. **最后合并：**丙方合并、按第五节分析、按第六节判定。

回溯版的两处弱点必须在发表时写明：其一，**事前变量是事后重建的**，即使严格设了数据截止日，重建者的注意力仍可能受档案的显著性影响；其二，事件与人员的入库不是前瞻确定的，存在幸存者偏倚（离职者的档案可能不全）。因此回溯版的结论**只能触发条款甲、乙、丙的“初步”版本**，须由一次前瞻版复核方可写入两章正文的定稿。

九、伦理条款（三条缺一不可，不得拆开引用）

其一，不得制造、放任或延迟事件。本研究的价值不得成为任何人容忍风险的理由。安全关键系统中，不得为“制造事件机会”安排新人顶岗或撤走冗余；任何执行方若发现响应安排因本研究而改变，须立即中止并申报。这一条与《峰兑》伦理三条之②、本书第十七章第二条同源。

其二，读数指位置与制度，不指人。低 β 、低 S 的人首先记录的是**他被安排在什么位置上**——有没有让他亲手做事的机会，有没有让他到场的事件。本研究的结果不得用于裁员、降职、调岗或任何不利人事安排；参与须自愿，且组织须提供对价（如本人可携带的兑付记录副本）。这一条与本书第十七章第一条同源，且是本研究向参与者的承诺，不是研究者的自我约束条款。

其三，编码规则公开、复核通道存在。三成分判据、兑付记录四条构成、结局三维度的全部编码细则，须随结果一并公开；任何被编码者可申请复核自己的编码，并可要求在发表数据中撤出（撤出数须报告）。

去标识与知情：结局评定使用去标识档案；参与者须在事前变量编码开始前获知本研究的存在、变量与用途。医疗场景另须遵从当地伦理审查程序，本协议不替代任何一处伦理审批。

十、既有印象申报

按《峰兑》第十四节的做法，协议撰写方须在取数前申报自己对结果的既有印象，以便读者判断判负条款是否被写得偏心：

本产线的既有倾向是 H1（《底流》侧）更可能成立——理由是本产线在写作《底流》时接触到的语言流失与技能保持文献普遍支持“停用即衰退”。为对冲这一倾向，本协议做了三件事：**第一**，条款甲（对《底流》不利）的触发只需效应不存在，不要求方向相反；**第二**，灵敏度内对照借自《峰兑》的娘家学科，其未达门槛时优先保护《峰兑》一侧（判未获裁决而非判其为伪）；**第三**，条款丁（两说各管一段）与条款丙（两说共负）明确写在表内，防止把结果硬读成非此即彼。

一处必须写下的自我限制：本协议对《底流》而言是“自己写的检验”，其独立性弱于外部提出的检验。若有任何一方愿意重写本协议的判负条款并执行，本产线接受其版本优先——包括接受更严的门槛。

十一、这场实验裁完之后

无论结果如何，都有三件事随之确定，写在这里免得事后追加解释：

其一，本书第十六章的**第 5 场与第 6 场同时结清**。该表两行合并为一行，两章各自的 F1 一并划账。全书十一场欠条自此减为十场——这是全书第一次真的还账。

其二，**第二章的两个部件各自受检**。H1 若被支持，存续条款（发生差要靠持续供给喂养）获得第一份真实数据；H2 若被支持，可读条件之二（事件处）获得第一份真实数据。若丙条触发（两说共负），第二章须删去存续条款与事件处条件各自的强表述，方程瘦两圈——按第三章失效条件二，这计入“判负累计”的两章。

其三，**无论哪一条触发，本次实验本身将成为一条兑付记录**——它有真实对手方（结果不由协议作者支配）、真实赌注（两章可能被判负）、无重来按钮（判负条款先于数据写死）、对手方签收（执行方公布结果）。按《底流》自己的三成分判据，这将是这条产线迄今第二次三成分齐备的执行。第一次是《底流》正文里那次自我判负。

十二、协议登记

- 本文件为 v1.0，写于任何数据接触之前。
- 执行方须在取数前将本文件全文与其 SHA-256 摘要一并存档（公开登记平台或组织内可公证的档案系统均可），并在发表时公布摘要。
- 存档之后，本协议不得修改。确需改动者，另立 v2.0 并同时公布 v1.0 与改动理由；v1.0 的判负条款不因 v2.0 的存在而作废。
- 本协议的哈希、功效脚本、编码细则与全部剔除记录，构成可复算路线；**该路线上没有一处需要协议作者在场。**

SHA-256（本文件正文段，自开头至“十二、协议登记”末条止）：

16e6c87277fd4de9deb144aaac8f91ad18288b126864e8c190f304e64b050339

（复算方式：取本文件自开头至上一条条目末尾的全部文本，UTF-8 编码后取 SHA-256。）

附：一页执行清单

1. 选场景，确认事件定义与档案可得性。
2. 指定甲乙丙三角色，做角色隔离，写下谁不能看什么。
3. 甲方封存能力清单，编码 d_{rel} 、 β 、 S ；**在事件之前完成。**
4. 存档本协议与哈希；此后不得修改。
5. 等事件自然发生。**不制造、不放任、不延迟。**
6. 事件后由第三方去标识，乙方两人评分，算 ICC。
7. 到达停止规则 ($N \geq 80 / 24$ 个月 / 退出) 之前不合并数据。
8. 丙方合并，先跑灵敏度内对照，再跑主模型。
9. 按第六节判定，照原样公布——包括对两章都不利的结果。
10. 把结果寄回给两章的作者，以及任何引用过这两章的人。

封底

当产物不再是证据，还剩下什么能把人分开？

生成式系统把一切可展示的产出抬到同一地板之后，简历、作品集、笔试与面试同时失去分辨力——不是候选人变得相同，是量具在灵敏度意义上死亡，而它照常出分。

本书收录十项独立研究：十组互不相识的学科三元组，各自推翻自己脚下的共同前提，最后汇到同一句话——**人工智能时代的核心竞争力，等于发生学能力的差异度，在其可读条件之内**。十条小径踩出同一个缺口，这个缺口没有人设计。

■ 十个可操作的读数：门位次、重走落差、留种率、脱稿成立率、在喂率、归读率、盲读率、未见率、他择率、记账方判据

■ 五条可读条件：跨配组、事件处、相对世界样本、他择时刻、未入账段

■ 四十五对两两分界矩阵，十一场写死判负条款的未跑实验

■ 两场对本书自身的预注册审计——其中一场把本书判负，照原样登载

■ 附录四：一份可直接执行的联合判决实验协议，作者本人被明文排除在裁决之外

书中两场自我审计，一场过关，一场把本书一处强主张判负；判负原样留在正文里，缺口由第十五章逐条补写并标明来源。一部主张“发生差要靠被检验的执行来喂养”的书，先得让自己被检验一次。

适合：用人单位与人力资源、高校教师与培养方、面临就业的毕业生、教育与劳动经济研究者，以及任何一个想知道自己还剩下什么东西不能被代做的人。

下一件你交出去的东西——

答案是什么时候到的？时刻是谁选的？记在谁的账上？

德麦国际出版社 DEMAI INTERNATIONAL PRESS · SDE Universes
sdeuniverses.com

德麦国际出版社 · SDE 学派专著第 103 号 ISBN 978-981-18-9243-1 US
\$23.40

版本与字数（机检实数） 成书：2026 年 9 月 3 日。德麦国际出版社 · SDE 学派专著第 106 号，ISBN 979-8-90690-041-8，定价 US \$23.40。全书七部十七章，另附前言、导读、导论、结语、参考文献与四附录；全文汉字 254,137（含本页），总字符（含标点）419,398；其中十章原文汉字 176,547，新写部分汉字 77,590。两场审计预注册 SHA-256：48a34a45723cf56923044599298ca284f817040ecdffa73c3c94e33c2e08904be；附录四协议正文段 SHA-256：16e6c87277fd4de9deb144aaac8f91ad18288b126864e8c190f304e64b050339。第四章语料截至 2026 年 9 月 3 日，复检周期十二个月。附录四所载实验未执行；协议进书不减少第十六章任何一张欠条。本书未经独立评审；分数留给不是本产线的手。