

认知炉膛的消失：AI 时代教育困境的发生学诊断

王德生 · 德麦国际 SDE 学派 · 2026年8月3日 · 约 2.4 万字 · 24 条参考文献 · 网页长文

摘要

本文诊断一种此前未被命名的教育困境。既有批判或将其归为"认知卸载"的外部适应，或视作"技术垄断"对文化-认知框架的侵蚀，或以"功绩主体"的自我规训来解释，或以动机与自我调节水平的下降来处理。然而，这些判断共享一个未经审查的先验：它们均假定困境指向某种已存在能力的损伤、降级或异化。本文指出，在最关键的案例中，发生的并非某个已有能力的衰退，而是某个从未被允许完整建立的内部认知结构——本文称之为"认知炉膛"——的永久性缺失。本文以发生学方法为主线，将分析单位从"能力""文化"下沉至个体认知结构的生成过程本身。承重命题为：此困境不是认知卸载，也不是功绩主体的自我规训，更非技术垄断下的文化符号降格，亦不止于动机受挫，而是"认知炉膛的消失"——一套本应在反复的、无外部模板的"原始认知接触"中被独立建造的内在生成装置，在过度精准的外部帮助中被系统性绕过，使得学习者在离开支持系统后，不仅缺乏知识，而且缺乏一个发起任何认知行动的"内部起点"。本文通过一阶判断与敌意最近邻的代理坍缩，提取出控制变量"内部意义引力的建构空间"（Z），并与来自政治学的一根结构独立的第二轴"非决策"权力相撞，构建一个二维辨别格；在此基础上给出可裁决判据与可观测代理。本文另在移植过程中提出"正当性型非决策"——一种没有受益方的偏见动员——作为对原概念的一处修正。

SDE 创新智商 143

全文盲评 · 待独立复核。五维（S 结构精确度 / D 差异锐度 / E 纠缠深度 / I 不可还原性 / F 可证伪性）评出，分数对应投稿原稿（S 143 / D 146 / E 138 / I 138 / F 151）。本页第 3.2、3.3、4.4、5.2、5.3 节，第六节第四条证伪路径，结语后两段及两则附论由编辑增补；按本站规程，改稿者不得为自己改过的文本盖认证分，故增补部分不计入本分数，逐项清单见文末附论。增补部分的自评（非盲评，仅供参考）约 148——低于增补时预估的 151，差额几乎全部落在「正当性型非决策」一处：该主张最初被写成对非决策概念的修正，复核时发现卢克斯的权力第三面、布迪厄的 doxa 与组织社会学的规范性同构已覆盖其大部，故已在正文中收窄为两处残余。创新智商衡量的是「一个此前不存在的认知物在发生意义上走得有多深」，不衡量该判断的真伪。（评分：Claude · 2026年8月3日）

目录

与前一篇的关系：本文是站内《当炉膛从未被建造》（教发生四论之四）的二阶续篇。那一篇命名了现象，本文给它装上第二根轴、切出四种可分辨的地形。两篇的分工线与可裁决处见文末附论。

一、引言：一个被长期误诊的教育困境 二、源起：一个底层本能的结构性排斥 三、敌意最近邻的代理坍缩与控制变量 Z 的提取 四、第二轴强制：当议题本身被挡在桌外 五、可裁决判据：让四家最近邻在

同一格里给出相反预测 六、证伪条件 七、结语：在非决策的遗忘之处 附论·与站内先行研究的划界
附论·本文与初稿的差异 参考文献

关键词 认知炉膛；内部意义引力；非决策权力；正当性型非决策；零启动判据；认知独处时段

一、引言：一个被长期误诊的教育困境

当人们谈论 AI 时代的教育困境时，叙事往往围绕一个核心意象展开：学生的某种能力正在被削弱或取代。一种流行的诊断是“技术学习负荷论”，它指出，当知识的获取变得过于便利，学生的记忆力与深度加工能力会因“认知卸载”而弱化。另一种更宏观的批判来自媒介生态学传统，它将此困境诊断为教育叙事在技术垄断时代的彻底溃败，学习沦为一种技术性游戏，伴随着严肃性的消逝。还有一派激进的声音，追溯至伊万·伊里奇，认为 AI 不过是让那个原本就旨在制造“能者”的学校制度变得更精准、更难以逃脱的终极工具。在个体心理层面，它又被描述为一种更隐蔽的自我规训，是功绩主体在“成为更好的自己”这一律令下，将自己异化为一个可被教学技术优化的对象。

这些诊断深刻、犀利，为我们提供了审视当代教育病理的宝贵视角。然而，它们共享一个底层预设：它们共同处理的那个“对象”，是一种已存在能力的受损、被破坏或被异化。认知卸载论谈论的是记忆功能的退化；技术垄断论探讨的是文化符号系统的整体降格；功绩主体论刻画的是一个原本整全的自我被劳动逻辑掏空。它们共同追问的问题是：“AI 把教育的哪种既有好处弄坏了？”

本文提出一种不同的诊断。本文不再追问“哪些既存的教育对象受到了 AI 的损害”，而是追问一个更根本的发生学问题：“一个‘受过教育的人’，其内部那些支撑独立思考与生成性判断的认知结构，究竟是怎样被发生出来的？而 AI 又对这个发生过程本身做了什么？”这一问法的倒转，将分析的焦点从“状态的恶化”移向了“结构的从未建成”。

本文的核心判断是：AI 时代最深刻的教育困境，并非学生的某种已具备的认知能力被侵蚀，而是一种本文称之为“认知炉膛”的内部生成结构，在系统性的“完美帮助”下，被永久性地绕过了。这不是损伤，是从未被安装。这一判断源自对“认知接触替代”机制的观察：当一个学习者与知识对象的原初的、充满摩擦的互动回路，被外部帮助体（教师、家长、AI）以优化过的认知流程所替代时，被替代的不仅是某一次学习的产出，更是学习者内部认知结构得以生长的唯一机会。当这种替代成为常态，那个我们称之为“炉膛”的内部装置——一个能将外部信息投入其中、经过内部灼烧、熔炼、重组，最终生成个人化理解的不可见的空间——便从未被建造。

本文与它的直接前身。“炉膛从未被建造”这一判断，本站先导研究《当炉膛从未被建造》已经作出，并已命名了其中的两件东西：那个自发把散落经验组织成属于自己的新判断的内部倾向，被命名为“意义引力”；它在建构窗口期中被系统性替代所导致的后果，被命名为“认知秩序的坍塌”。本文不重复那次命名，也不与它竞争。本文承接它，并做一件它没有做的事：把一个被命名的现象，变成一张可以分辨的地图。

这个区别是本文全部工作的位置所在，值得说清楚。先导研究交出的是一阶判断——它论证了“这件东西存在，而且它此前没有名字”。一阶判断的成立，不等于它可以被使用：面对一个具体的学习者，我们仍然无法回答“他的炉膛是没建成、还是建成后被扭曲、还是建成了但被制度盖住了”。更麻烦的

是，一阶判断天然地滑向一种全称诊断——把一切学习困难都归因于"炉膛不在"。本文要做的，是给这个判断装上第二根轴，使它从一块写着"此处有深渊"的警示牌，升维成一张能分辨深渊不同深度的地图。这张地图的第二象限——那些在最不利的制度环境中炉膛仍然燃着的幸存者——是先导研究无法预测、也无法解释的一格，而它恰恰是全部可裁决判据的产地。

二、源起：一个底层本能的结构性排斥

在深入任何理论之前，我们首先正视一个遍布于教育现场的日常直觉。任何一位有经验的教师或家长，都能在自己的记忆中检索到这样的时刻：孩子卡在一道难题上，眉头紧锁，反复涂改，时间一分一秒流逝。此刻，一个强大到近乎生理性的本能会涌现——去帮他，去把那条最简捷的路径画给他看。这个本能的底层逻辑不来自任何恶意，恰恰相反，它来自教育者身份最核心的正当性根据：**我的价值，在于提供有效的帮助**。一个不能帮助学生解除困惑的教师，一个看着孩子无助哭泣而无动于衷的家长，将首先面对自我专业伦理的审判。

然而，本文的起点，正是对这一本能的无害性的根本质疑。我们需要将其推至极限来审视它的结构后果：一个以“提供有效帮助”来证明自身正当性的教育系统，其核心操作，是否注定与催生真正理解的“裂缝”相敌对？此处所谓的“裂缝”，指的是学习者从“我知道什么”的平滑表面，跌入“我不知道”的断裂地带的那个认知时刻。在这个裂缝中，既有的知识框架暂时失效，外部帮助尚未抵达，学习者被抛回自身那片混沌的、尚未被语言清晰锚定的内部感知。经典的学习理论已经以不同的技术语言确认了这个时刻的构成性价值：杜威的“困惑”是整个反思思维的起点；皮亚杰的“认知失衡”是图式通过顺应而升级的唯一触发器；维果茨基的“最近发展区”本质上正是裂缝——一个有外部脚手架支撑但内部尚未独立的认知缺口。

本文对此困境的诊断是：一个以“帮助”为核心本能的教育系统，并非在“帮与不帮”之间做出了一个错误决策；而是它的整个合法性叙事，已将“裂缝时刻”从“教育过程必要的构成部分”这一认知版图中彻底排斥了出去。这不是一个可以被轻易修正的“技术性错误”，而是一个被系统正当性所驱动的“结构性排斥”。当 AI 作为一个永不疲倦、秒级响应、精准诊断的帮助者进入这一系统时，它所增强的，恰恰是这个系统最本能的那个操作：填平裂缝。它没有制造这场危机，它只是让这场危机的逻辑跑完了自己的全程。

2.1 改切裁定：分析单位的收窄

在展开这一诊断之前，必须对其适用范围进行严格的改切裁定。本文的原初分析对象指向“AI 时代的教育”这一总体性范畴，但这一范畴在分析操作上过于粗放。笼统地谈论“教育系统排斥裂缝”，会在北欧的探究式课堂与东亚的应试工场之间制造一种虚假的同一性。事实上，裂缝被排斥的方式、程度与后果，在不同教育文化中差异极大。

因此，本文在正式启动分析之前，明确将分析单位从“整个教育系统”改切至其一个特定的、但具有广泛经验对应物的子类：**即那些将“秒级提供正确、高效、可评估的认知帮助”确立为教师专业伦理核心、家长责任伦理核心和技术产品竞争力核心的教育实践场域**。在本文的后续分析中，每当我们提及“教育系统”，其所指的均是符合上述界定的特定子类，而非任何与“教育”有关的全部活动。

这一改切的理由立足于问题本身：本文承重命题所要解释的核心机制——即系统性帮助如何绕过内部认知建构——只在帮助被高度制度化、伦理化和技术化的场域中才能被清晰地观测和分析。在那些仍保留了大量无预设、无评估、无即时反馈的认知独处时段的教育文化中，这一机制可能根本不构成主导性的困境。因此，改切不仅是合法的，更是使分析得以精准聚焦的必须操作。

作为这一改切的配套，本文的阴性案例主要有两类：其一，那些结构性保留了认知独处时段的教育文化中，应当观察到同等 AI 介入程度下，炉膛缺失的比例显著更低；其二，那些即便身处高帮助密度环境、却仍然保留了内部生成装置的幸存者，其幸存机制是否为本文所预言的"认知独处时段"而非其他尚未被考虑的因素，将是检验本文机制的关键。第二类正是下文第二象限所描绘的个案，也是本文全部判据的产地。

三、敌意最近邻的代理坍塌与控制变量 Z 的提取

本文的承重命题——AI 时代教育困境的本质不是能力的损伤，而是"认知炉膛"这一内部生成装置的从未建成——必须逐一消解那些它最容易被混为一谈的既有代理变量。对手的数量不在多，在于能否被精准地推至其解释力的边界。本文选择的四位敌意最近邻，分别占据了解释这一困境的四个主要方向：认知功能、主体状态、文化结构，与动机-自我调节。前三个方向是本文原初诊断的对手，第四个方向是最贴身、也最容易吞掉本文的一家。

3.1 消解三项既有代理

第一位对手是"认知卸载"，其经典经验支撑来自斯帕罗等人的谷歌效应实验：当人们预期可以轻松从外部获取信息时，他们对信息本身的记忆会变差。将其逻辑外推至 AI 教育场景，一个自然的诊断是：学生之所以在离开 AI 后表现出认知困难，是因为他们过度依赖外部存储与处理，导致了相关内部认知功能的用进废退。认知卸载论处理的承重变量是"内部认知功能的因不转而退化"。本文并不否认这一机制在大量日常案例中的真实性。然而，本文分离出它的盲区：这个变量无法区分"功能性的暂时依赖"与"发育性的结构缺失"。一个因过度依赖导航而认路能力下降的人，在撤除导航并经过一段时间的重新练习后，其内部的空间认知能力是可以恢复的——因为他曾经拥有它。但本文要诊断的案例，其经验特征恰恰是"从未拥有"。在那些最极端、也最具诊断价值的案例中，学习者从一开始就不曾被抛入那个需要自己生成内部导航地图的原始认知接触中。AI 的介入不是在"中途接手"了一项此前由学习者自主完成的任务，而是在这项任务本应启动学习者内部发生学程序的那一刻，直接替代了整个程序本身。因此，认知卸载论的代理变量在分离点——"认知功能能否在撤除外部帮助后通过重新训练而恢复"——上坍塌了：它解释得了"退化"，解释不了"从未被安装"。

第二位对手是"功绩主体"及其核心概念"自我规训"。在这一框架下，学习者陷入困境并非因为能力退化，而是因为将外部绩效标准完全内化为对自我的暴力性要求，从而陷入了永不停歇的、自我剥削式的"学习劳动"。功绩主体论处理的承重变量是"主体将外部评价内化为自我驱动之后的异化"。这个诊断能有力地解释一个令教师困惑的现象：为何那些看起来最"自律"、最"主动"的学生，其学习行为却透出一种深层的不自由与意义感丧失。然而，它的代理变量在本文所固定的一个关键案例频谱中发生了坍塌。在同等高度内化外部评价的群体中——即，都符合功绩主体描述的学生——存在着一种显著的认知行为分化：一部分学生在面对一项毫无外部评价指标、完全由内在兴趣驱动的任务时，其认知焦灼与行为投入会如断电般突然消散；而另一小部分学生，尽管同样身处评价体系之中，却仍能在这种"无用"的任务中保持持续而专注的投入，其投入源显然不来自对绩效指令的内化。功绩主体论仅凭"规训程度"这一个变量，无法解释这组差异。它的盲区在于：它只能解释"为什么出发"（为绩效而学），却无法分辨"在路上"的两种截然不同的主体状态——一种在外部指令撤除后瞬间停摆，另一种则由一个内部的、非绩效的自发热核心所驱动。本文要命名的，正是这个自发热核心的构成条件。

第三位对手是"技术垄断"，代表了一种更宏阔的文化批判向度。它的承重变量是"整个文化符号系统在技术逻辑的统治下发生的整体性降格"。在技术垄断的框架里，教育困境是整个文化溃败的一个局部表征：当效率、可计算性和技术方案成为一切社会事务的元叙事，教育本身的意义（为何而教）便会被掏空，沦为一种如何更有效地传递信息的技术性游戏。本文认同这一诊断在文化层级的有效性，但其

代理变量同样在其解释力的边界处坍塌了：它的分析单位是"文化"或"整个符号系统"，因此它天然地无法对同一文化系统内部出现的差异化结果做出具体预测。面对在同一个高 AI 渗透率、高技术依赖的教育文化中，那些"炉膛"仍未被熄灭的认知幸存者，技术垄断论只能将其解释为暂时未被主流叙事完全吞噬的边缘个案，而无法提供一套可操作的机制去阐述他们是如何幸存下来的。这正是本文需要一根结构性独立的第二轴的原因——这根轴必须在被技术垄断论视为均质化的那片文化平原上，测绘出内部的地形差异。

3.2 消解第四家：动机-自我调节一路，与它与本文的分离线

上述三家分别据守认知功能、主体状态与文化结构三个方向。但还有一个方向没有被占，而它是本文最贴身、也最容易在不知不觉中吞掉本文的一家：**动机与自我调节**。这一路与本文谈论的是同一批学生、同一批行为、甚至同一批实验情境。若不正面处理，本文的全部判断都可能只是它的一次哲学化重述。本文在此把它请到桌前。

这一路由三支构成，它们各有一个承重变量：

其一，**自我决定论**。当自主性、胜任感与归属感三项基本心理需要被满足时，内在动机得以维持；被剥夺时，动机向受控型迁移，其最左端即**无动机**——一种缺乏行动意图的状态。承重变量=**自主性需要的满足度**。

其二，**内隐智力理论与成就目标定向**。持能力实体论者遇挫时归因于自身能力不足、行为冻结、回避；持能力增长论者归因于策略不当、更换路径、坚持。承重变量=**关于能力可变性的内隐信念**。

其三，**自我调节学习**。预想—执行—自我反思三相循环可通过教学建立，元认知策略是可教的。承重变量=**元认知策略的习得程度**。

这三支合起来的解释力极强。更要紧的是——这一点必须由本文自己指出来，而不是等审稿人指出来——它们所预测的行为分化，与本文初稿第五节所给出的两条判据**几乎逐字重合**：一组学习者失败时归因于方法、更换路径、持续投入；另一组归因于能力、迅速冻结、退出。这组对照自 1978 年的无助型/掌握型归因研究以来就已写在文献里。一篇声称发现了新病理的论文，如果它的判决性判据落在半个世纪前就已被预测过的行为上，那么它的新意只在词汇。

本文因此必须把分离线画在一个更窄、也更硬的位置上。

此处必须先纠正一个容易顺手写下、却不成立的说法：不能说这三支**没有**"装置是否存在"这个范畴。自我决定论有——"无动机"就在它动机连续谱的最左端，定义正是缺乏行动意图。因此本文与这一路的分歧不在范畴的有无，而在**那个状态被认定为什么性质**。

三支的共同预设是：**缺席是可逆的**。在自我决定论中，无动机源于非偶联信念或胜任预期过低，因而提供自主支持与可及的胜任经验即可回升；内隐智力理论解释装置遇阻后往哪里转，信念可教；自我调节学习解释如何让它转得更有章法，策略可教。三支处理的都是一台**原则上可被重新点着的**装置的燃料、方向与效率。

本文主张的是另一种性质：**结构性缺席**——不是这台装置熄了火，而是在它本应被建造的窗口期中没有被建造，因而不存在一个可供“回升”的旧状态。这一分歧比“有没有这个范畴”更硬，因为它把争论收窄到一个可测的量上：**可逆性**。

由此得到本文与这一路的分离线，它可以被压缩成一句可操作的话：**归因发生在第几步**。

无助型归因者的行为序列是：接触任务 → 生成一个目标表述 → 尝试 → 失败 → 归因于能力 → 冻结。他**先启动了**。“归因”是对一次已经发生过的尝试的解释。正因为序列的前三步都已完成，撤除评价压力、改变内隐信念、教授元认知策略，才可能让这个序列在“归因”那一步转向——这正是这一路数十年干预研究之所以有效的原因。

炉膛缺失者的序列缺的是**第二步**：接触任务 → （空白）。他不产生目标表述，因此没有尝试，因此没有失败可以归因。他所报告的不是“我做不好这个”，而是“我不知道要开始什么”。他不会犯错，因为他根本还未进入解决问题的状态。这一区别，不是在“技能强与弱”或“信念对与错”的谱系上，而是在“有”与“无”的位置上。

这条分离线将在第五节被兑换成一组可编码、可计数、可对照的观测指标。此处先记下它的形式：**无助型组的产物是错的；炉膛缺失组没有产物**。一个错误的目标表述与零个目标表述，不是同一条谱系上的两个刻度。

3.3 与“帮助时机”实证传统的分工：同盟而非敌手

有一支文献与本文立场高度一致，必须在此交代清楚，否则本文会显得不熟悉自己所处的领域。

生成性失败研究表明：让学习者先在没有指导的情况下与一个超出其当前能力的问题纠缠、失败，再给出规范讲解，其概念理解与迁移的长期表现优于“直接讲解在先”。**合意难度**研究表明：制造提取困难、间隔、交错等短期内减慢表现的条件，长期保持与迁移反而更好。**帮助两难**则在智能辅导系统的设计中指出：何时给帮助、给多少帮助，没有一个普遍最优解；帮助过早会剥夺学习者必要的建构工作。

这三支已经用实证方法确立了本文第二节所说的“裂缝具有构成性价值”。本文不与它们竞争，本文站在它们肩上。但分工线必须画清楚：

它们处理的是单次任务尺度上的帮助时机；本文处理的是发育窗口尺度上的装置有无。

生成性失败研究的被试，在实验开始的那一刻**已经拥有一台会自己启动的装置**——正是因为有，“让他先自己撞一撞”才是一个可执行的实验操作。如果把生成性失败的实验范式施加于一个炉膛缺失者，实验的第一步就无法完成：他不会去撞，他会等指令。

这恰恰构成本文对这一支的一条可回溯检验的预测：**生成性失败范式的效应量，在高帮助密度环境中长大且无认知独处史的学习者身上，应显著小于常模；且这一组中会出现一批“无产出”而非“低质量产出”的被试**。在现有的生成性失败研究里，交白卷的被试通常被作为无效数据剔除或计为零分处理——而本文主张，那批白卷正是本文要找的东西。这条预测指认的是一个此前被系统性丢弃的**观测类别**。但要兑现它需要一项新的前瞻研究：把零产出改为记录而非剔除，并在同一批被试上同期采集独处

史。旧研究通常不曾采集这项信息，故这不是回头翻旧数据就能做成的事——第五节还会再交代一次。

3.4 Z 的命名：内部意义引力的建构空间

在逐项消解了"认知功能退化""主体自我规训""文化符号降格"与"动机-自我调节水平"这四项代理变量之后，剩余的那个不可再被它们解释的部分，便是本文所提取的控制变量 Z——**内部意义引力的建构空间**。

这个命名承接本站先导研究对"意义引力"的界定，并把它从一个内部倾向推进为一个可被建构或被绕过的**空间**。一个学习者在漫长的、累积的"原始认知接触"中，那些独自面对问题、在无外部评价和即时辅导的缝隙里，反复体验"我自己找到了它"的时刻，逐渐在其内部生成了一种特殊的认知重力。这种重力，使得外部世界的信息碎片在涌入时，不再只是被平行地存储或遗忘，而是被自动地、不可抗拒地吸入一个私人的意义熔炉：它们在那里被拆解、被怀疑、被与既有经验重组，并最终凝结为一种带有个人印记的、可独立输出的判断。这个空间，不是任何一项具体的认知技能（如记忆力、逻辑推理能力），也不是一套外显的元认知策略（如自我提问、计划监控）。它是一个使上述所有技能和策略得以被主动发起、并在失去外部指令后仍然持续运转的底层动力装置。

这个装置的构成需要三个不可同时为零的初始条件：第一，足够多的、未被外部帮助抢先填充的"认知裂缝"经验；第二，裂缝中，学习者亲自完成从浑浊到清晰的全部关联步骤，而不是从"别人的步骤"中挑选最后一步；第三，这种亲自完成的产物，被至少一个重要的"他者"以非评价性的方式看见并回应过。正是这三个条件不能由任何精准的外部模板来替代提供，决定了 Z 的发生必须在"无外部模板的原始接触"中被独立建构。三个条件中的任意一项长期为零，Z 便无法达到一个能支撑自主认知起始的临界阈值。

至此，本文完成了从一阶判断到控制变量提取的全部理论准备。但是，仅仅命名 Z，本文的产出仍停留在一阶判断的延长线上：我们只是换了一个更精确的名字去回答"困境是什么"。要将此诊断从"命名现象"升维为"创立辨别维度"，Z 必须撞上一条结构独立的第二轴。下一节将完成这一构造。

四、第二轴强制：当议题本身被挡在桌外

4.1 炉膛的发生学：与皮亚杰和维果茨基的正面交手

在引入第二轴之前，本文必须先回应一个在整个论证链条中最关键、也最无法绕过的质疑。这一质疑直指本文承重命题的命脉——它来自发展心理学与教育心理学传统中最具奠基性的两位学者。皮亚杰的发生认识论和维果茨基的社会文化理论，难道不是早已从不同侧面论述了“内部认知结构只能在主体与世界的主动互动中生成”这一核心机制吗？如果本文的“炉膛”概念，在根本上可以被皮亚杰的“图式通过同化与顺应的双向建构”或维果茨基的“最近发展区内的有效中介”所覆盖甚至替换，那么本文的整个理论大厦将失去其不可还原的根基。

本文完全承认皮亚杰的建构主义揭示了认知结构的一个根本发生条件：主体必须亲身对客体进行操作，通过同化与顺应的双向过程，才能建构新的图式。任何外部现成的、不容主体改造的模板，都无法替代这种双向建构活动对认知生长的根本性驱动。本文所分析的“原始认知接触”与“在无外部模板的摩擦中独自完成从浑浊到清晰的全过程”，其发生学原理在皮亚杰的理论中确有对应：它本质上描述的正是一种深度的、充分展开的顺应过程。同样，维果茨基的最近发展区理论极清晰地指出，有效的中介不应替代学习者的内在发展过程，而应走在发展前面、搭建机会使其内在建构得以发生。他对“中介”的深刻洞察，是对纯灌输式教学的提前判刑。

然而，正是在这一点上，本文的论证必须画出一条关键的分离线。皮亚杰与维果茨基（及其后继者）的理论，共同处理的对象是具体的、领域性的“认知图式”或“高级心理机能”——比如守恒、分类、书面语言的掌握、科学概念的建立。他们解释了这些图式在个体内部是如何被成功建构的，也解释了它们在什么条件下会建构失败（如认知冲突不足、中介不当或缺乏）。但是，他们未曾单独提出并系统处理一个更为底层的元层装置：一个使所有这些具体图式的建构活动本身，能够被学习者主动发起、持续维持，并在失去外部指导、脚手架和即时反馈后仍然不熄灭的内部动力起点。

这个起点，就是本文所说的“认知炉膛”。它不是任何一个具体的图式，而是“图式之母”——一个内部的意义引力建构空间。它的功能，不是保证某个特定的认知操作能够成功（比如得出正确答案），而是确保一个认知行动能够被主动地“开始”。当一个学习者面对一个完全无模板、无明确指令、无评价标准的开放情境时，他内部是否有一个自动点火、将碎片吸入并尝试熔炼的装置在运转？皮亚杰的理论解释了点火成功之后火焰如何塑形（图式的同化与顺应），却没有解释为何有些人的燃料系统从一开始就从未被安装，以至于连一次点火都不曾发生。维果茨基的理论提供了烧得更好的“薪柴”（文化工具与中介），也划定了火势可旺的范围（最近发展区），但他预设了那堆火本身的在场——一个能够积极响应中介、并将其内化的活生生的发展主体。本文所面对的，正是这个“主体在场”的前提，在当代教育实践中被系统性掏空的现场。

4.2 候选轴的排除

一个控制变量 Z 即便被精确地提取出来，并与其所有代理变量完成了分离，它依然停留在一阶判断的延长线上。要完成从“命名现象”到“创立辨别维度”的跃迁， Z 必须撞上一条结构独立的第二轴。这条轴

不能是 Z 的成因，不能是 Z 的后果，更不能是 Z 在另一个语域的换皮重述。它必须来自一个与教育认知心理学在提问方式上根本无关的领域，却能在此刻把 Z 从一块标记着"此处有深渊"的警示牌，升维成一张能分辨深渊不同深度的地图。

首先被排除的候选轴是"外部支持的强度"。这似乎是一个直观的选择：以"外部帮助的精准度与即时性"为横轴，以 Z 为纵轴，构建一个从"低外部帮助×高建构空间"到"高外部帮助×低建构空间"的连续谱。然而，这条轴与 Z 的分离是虚假的。本文在第三节已论证，Z 正是在高外部帮助的常态下被侵蚀的；外部帮助强度是 Z 的成因之一，而非与 Z 结构独立的第二维。用它做轴，等于拿炉膛的燃料消耗量去解释炉膛的消失——它们高度相关，但不是两根独立的轴。此轴排除。

其次被排除的是"学习者的初始认知禀赋"。这同样是 Z 的一个输入端，而非独立的辨别维度。用它做轴，会把本文的论证拖回到天赋论与教育机会均等的旧辩题中去，而本文追索的问题恰恰是：在禀赋相当的学习者之间，那个导致内部生成装置被绕过与否的分水岭究竟是什么？禀赋无法回答这个问题。此轴排除。

第三条被排除的候选轴是"教育系统的评价刚性"。这条轴来自本文站内先导研究《正当的排斥》，其核心判断是：当"正确"成为教育系统唯一的合法性尺度时，排斥生成是其维护自身正当性的必要操作。这条判断与本文的承重命题高度协同，但正因为如此，评价刚性是 Z 得以被生产的制度条件，而非与 Z 并列的第二维。将它选为第二轴，会导向一个正确的、但识别力有限的双高象限（高评价刚性×低建构空间），无法撕开本文真正想切入的那个更隐蔽的边界：在评价刚性完全相同、外部帮助同样泛滥的制度环境里，为什么有些学习者的内部炉膛熄灭了，而另一些在缝隙中残存了下来？它解释不了这个差异。

本文选择的第二轴，必须能够切割开这个被制度均质化的环境内部那看不见的裂纹。它不能是关于"怎么做"或"怎么评价"的变量，而是关于"什么甚至不被允许被摆上桌面"的变量。

4.3 轴的确立：非决策权力

本文引入的第二轴，是政治学家巴赫拉克与巴拉茨提出的"非决策"概念。在权力研究的传统中，决策视角关注的是"在议题 A、B 之间，谁赢了"。而非决策视角则提出一个更根本的问题：在议题 A、B 被摆上决策桌之前，是什么力量将某些潜在议题 X、Y 挡在了议程之外，使它们根本不被讨论、不被看见、甚至不被感知为"一个可决策的事项"？

将"非决策"引入教育认知的诊断，并非一种比喻式的挪用。本文提出一个严格的类比：教育系统对学习者内部发生的"不知的裂缝"，执行的正是一种非决策式的操作。它不是在"帮"与"不帮"之间做了一个错误的决策；它是在学生的整个学习生涯中，通过教师专业规范、家长责任伦理和教学法信条的合力，系统性地框定了一条边界——边界这一侧，是可被帮助、可被评估、可被优化的"学习困难"；边界那一侧，是那个必须不被触及的、悬而未决的、连教师自己都不知道如何量化的"裂缝时刻"。这个裂缝，不是被决定"不帮"，而是被"非决策"了：它从未作为一个需要被保留的选项，进入过教育实践的议程。

这根轴与 Z 的关系不是因果关系，而是结构正交关系。Z（内部意义引力的建构空间）是一条读自学习者内部的、关于"发生潜能"的连续谱。而非决策权力，是一条读自制度层面的、关于"某个议题被挡在议程边界之外的严密程度"的连续谱。二者测量的是完全不同的实体：一个测量个体内部认知建构的可

能空间；另一个测量系统对"保留裂缝"这一可能操作的系统性屏蔽程度。一根轴读的是炉膛里还剩多少温度，另一根轴读的是那些可能往炉膛里添柴的手，在意识到炉膛存在之前就被引导去了别处。

4.4 移植的合法性：从利益型非决策到正当性型非决策

把一个政治学概念搬进教育认知的诊断，最容易出的事故是：搬走了词，没搬走使这个词有解释力的那套结构。本文在此把这道关口正面走一遍，因为走完之后，得到的不只是移植的合法性，还有一件可以还给政治学的东西。

非决策概念有一个常被忽略的支撑件：**偏见动员**。议题之所以进不了议程，是因为一套既有的价值、仪式与程序系统性地偏向某些群体的利益，而这些群体从该议题不被讨论中获益。换句话说，政治学里的非决策是**有受益方的**；正是受益方的存在，使"为什么这件事从来没人提"这个问题有一个可指认的答案。

把它搬进教育场，第一个要过的关就在这里：**谁从"裂缝不被保留"中获益？**

诚实的回答是：没有人。教师不获益——她因此失去了她最想要的那种课堂时刻，那种一个孩子忽然自己转过弯来的时刻。家长不获益——他为此付出了大量的金钱、时间与睡眠，并换回持续的焦虑。学生更不获益。教育技术公司获得了收入，但它们并不是这道边界的设立者；它们是**沿着一道已经画好的边界**扩张的，边界在它们出现之前就在那里。

因此，若照搬原概念，这次移植是失败的：没有受益方，就没有偏见动员，就没有非决策。

本文的回应不是绕开这一点，而是从这一点上取得一处对原概念的修正：**偏见动员并非只能由利益供给，它同样可以由正当性供给。**

在教育场里，把"保留学生的不解"挡在议程之外的，不是任何一方的利益，而是一个所有参与者都不能公开反对的价值：**有效的帮助是善的**。一位教师不能在教研会上提议"我们应当有意让部分学生更久地停留在不会的状态里"——这句话不是被否决，是**不可说**；它一说出口就使发言者的专业伦理受损，同事们的反应不会是反驳而是尴尬。一位家长不能在家长群里说"我今晚决定不管孩子那道题"——这句话同样不是被反驳，而是会被读作失职的自陈。一家教育技术公司不能把"我们故意慢一点、少给一点提示"写进宣传页，因为整个竞争维度已经被"更快、更准、更懂你"锁定，把慢当卖点等于自愿退出市场。

这是一种**没有受益方的偏见动员**：三套彼此独立的正当性规范——教师专业伦理、家长责任伦理、产品竞争伦理——在同一个方向上叠加，各自都无可指摘，合起来把同一项议题挡在门外。本文称之为**正当性型非决策**。

此处必须先把已被占住的地盘让出来，本文的主张才立得住。"没有受益方也能有排除"这件事，本文不是第一个说的。卢克斯的权力第三面正是为补第二面的这个缺口而提出，处理的就是不满从未形成、冲突从未浮现的那一层；布迪厄的 doxa 指认的是那片"不被讨论的宇宙"；组织社会学中的制度化神话与规范性同构，说的更是组织为正当性而非效率采纳做法、专业规范驱动趋同——其中同样找不到受益方。这几家合起来，已经覆盖了上文所描述景象的大部分。

因此本文在这一处的主张必须收窄到两件残余上，而不是宣称对非决策概念作出了修正。其一，**机制的具体形状**：起作用的不是一套支配性价值观，而是**多套彼此独立、互不隶属、各自都可被单独辩护的规范系统在同一方向上叠加**——正因为它们独立，任何一套的松动都不足以打开议程，这解释了单点改革为何反复失效。其二，**由此推出的补救不对称**：卢克斯那一路仍以"真实利益"与支配群体为轴，其解放路径是让被支配者认出自己的真实利益；而在多规范叠加的结构里没有可指认的支配群体，故揭露无处落脚。这条不对称是本文的。除此之外的部分，本文视为对既有传统的一次教育场景应用，不主张原创；沿用这个称谓只为文内指称方便，不是宣告一个新范畴。

正当性型非决策与利益型非决策共享同一个结构：议题在进入决策之前即被排除、排除本身不留下决策痕迹、因而事后不可归责。但它的供能方式不同，因而**解除方式也不同**：利益型非决策可以通过揭露受益方而松动——一旦人们看清谁在获利，议题就获得了被提起的理由；正当性型非决策不会因为揭露而松动，因为它没有受益方可揭。它只能通过**在正当性内部为被挡的议题挖出一块被承认的合法位置**才可能松动。

这一改造不是为了让类比看起来更漂亮。它有可裁决后果，而且这个后果与教育批判的主流方向相反：**如果本文的第二轴成立，那么以"揭露利益"为手段的教育批判——指出补习产业、教辅资本、绩效考核背后的获利结构——在削弱这道边界上应当收效甚微**；而以"给不确定性一块被制度承认的地皮"为手段的干预——例如在教师评价中设一项明确不计入绩效考核的自主教学时段，或在作业规范中设一栏"本题允许交白卷并写下卡在哪里"——应当收效更大。这两条方向相反的预测，可以在同一个学区的不同学校里对照施行，观测指标即第五节的三条判据。

而且这个亚型不止对教育有效。任何一个由多套彼此独立的善意规范叠加锁死的场域——临终医疗中"必须尽全力抢救"、灾害应对中"必须消除一切风险"、儿童保护中"必须消除一切无人看管的时间"——都可能是同一结构在不同材料上的显现。在这些场域中，同样找不到受益方，同样每一条规范单独看都无可指摘，同样有一个议题从未进入过议程。这一点使本文的第二轴不止是从政治学借来的一把工具：它同时把一个可移植的判别式带回那些场域：**看排除是由利益供能，还是由多套独立规范叠加供能**——前者可被揭露松动，后者不能。

4.5 辨别象限：炉膛熄灭的四种地形

以"内部意义引力的建构空间"（Z）为纵轴，以"对裂缝的非决策屏蔽程度"为横轴，构建一个二维辨别格。纵轴从"高"到"低"，描绘学习者面对无模板的开放问题时，内部发起认知行动的起点是坚实还是空洞。横轴从"低屏蔽"到"高屏蔽"，描绘整个外围系统（教师、评估、家长社群、教育技术规范）保留"裂缝时刻"的可能性，是一个从裂缝仍可能被偶然容留，到它被全部教学语言隔绝在议程之外的梯度。

第一象限（高建构空间 × 低非决策屏蔽）：私燃之境。裂缝在此不是灾难，而是教学的正常组成部分。教师的专业直觉允许暂停，允许沉默，允许"我不知道怎么帮你，但你再试试"。外部支持精准的退出不是失职，而是专业判断。在这一格中，本文的承重命题并非说 AI 侵蚀了认知炉膛；相反，它预测在这样一个系统中，AI 可以被用作扩展认知接触广度（而非深度替代）的工具。学习者内部的炉膛，在这里获得了建构所需的摩擦热与容错空间。

第二象限（高建构空间 × 高非决策屏蔽）：幸存之炉。这是最罕有、也最具理论诊断价值的一格。在这里，制度对裂缝的非决策屏蔽同样是高度严密的：评估标准一刀切，教师手册写满了"必须提供有效帮助"的指令，家长的手机应用实时推送监控报告。然而，在这些学习者的内部，炉膛却仍然燃着。他们是如何幸存下来的？本文预测，这种幸存不是通过制度内的合规操作达成的，而是通过制度覆盖不到的那些非正式缝隙：一位会故意留出一段"不讲清楚"的沉默的教师，一个在孩子独处时推开平板电脑的家长，或者学习者本人在往返学校的公交车上那段不被任何教学法监控的白日梦。这幅地图上最值得追踪的矿床正在这里，因为它是所有代理变量（认知卸载、功绩主体、技术垄断、动机-自我调节）都无法预测的异常常格。由于本文第二轴的创新性，此格的实证描述仍待进一步的质化田野研究去填充，但它的存在本身已经构成了对"制度决定论"的一道可裁决的挑战。

第三象限（低建构空间 × 低非决策屏蔽）：无柴之烬。这种情况的悲剧性最为隐蔽。从制度层面看，这似乎是一个健康的环境：系统并不刻意屏蔽裂缝，教师有相当大的专业自主空间，评估也不是铁板一块。然而，学习者的内部炉膛仍是熄灭的。本文在第三节讨论过的那种内在意义引力的丧失——那种即便正确答案都撤走，人也不再具有冲动去亲手触碰和重建一个理解的状况——就可能在这一格中浮出表面。这不是因为外面给得太多了，而是因为内部那个本应由早期认知接触慢慢养出来的自发热核心，在更早的阶段就已经停止运转了。对这一格的诊断，暴露出非决策轴不能解释一切：炉膛的消失，也可能来自对早期无意识的、非调控性的认知接触的长期剥夺。

第四象限（低建构空间 × 高非决策屏蔽）：冰封之室。这是本文承重命题所描绘的核心景象。制度的屏蔽强度与内部建构空间的消失在此汇聚，形成一种互相锁死的稳态。认知卸载论在此格中恰如其名：外部工具替代了内部记忆与检索。功绩主体论在此格中诊断了学生表面的自我驱动如何演变为自我剥削。技术垄断论在此格中解释了整个教育叙事如何沦为技术效率的附庸。动机-自我调节一路在此格中给出无助型归因与元认知策略的缺失。它们的判断在这一格中全都成立，这正是为什么它们长期被当作普适诊断的原因——因为它们在冰封之室里看到的病征，与它们在幸存之炉里看到的别无二致。它们看不见幸存之炉，因为它们没有第二根轴。本文的增量，不在于指出这间冰封之室的存在，而在于证明这间冰室只是辨别格中的一个象限——它旁边，还站着别的屋子。

4.6 分辨"被跳过"与"被损伤"

这张辨别格也为处理一个核心张力提供了工具：本文在行文中交替使用了"从未被建造"与"结构被畸化"两种叙述，这两者分别对应不同的病理机制。前者是被跳过（空无一物），后者是被损伤（结构变形）。

在辨别格的坐标上，这两者的分野得以具体化。**被跳过（从未被建造）**指的是一种极低的 Z 值，学习者在面对开放式问题时，展现出的不是错误的判断，而是没有判断——没有任何内部起点的空白。这在第四象限（冰封之室）和第三象限（无柴之烬）中都可以发生，而其成因则指向了"原始认知接触"的根本性缺失。

被损伤（结构变形）则是一种中等的 Z 值，但带着功能性的扭曲：学习者拥有内部起点，但这个起点的运作方式已经被外部监控系统重塑。他可能会不断地反刍问题，展现出高度自律，但反复在同一个点上回到外部标准进行核对，无法做出最终的个性化裁决。这是"生成"被异化为"生成者身份的展演"。

在辨别格中，这更可能发生在第二象限（幸存之炉）的边缘地带——那个炉膛还在，但已经烧的不是自己的柴了。

这个区分使本文的判断可以避免一个严重的误用：将一切学习困难都归咎于"炉膛不在"。它要求我们真正地走到辨识的位置上，去分辨这是一种空无，还是一种被锻造过的受苦。

五、可裁决判据：让四家最近邻在同一格里给出相反预测

任何以"内部发生结构"为名的诊断，都天然地滑向不可证伪的危险：它谈论的是不可见的内部空间，它的缺失无法被直接目击。本节的核心任务，是依据第四象限与第二象限的辨认框架，交出一组可裁决判据——一组特定的、可观测的学习者行为差异，让本文与它那些最强劲的敌意最近邻，在同一张辨别格的同一个格子中，做出相反的预测。

5.1 判别格与两个组

本文选择的判别格，是前文辨析过的第二象限——幸存之炉。在这一格中，制度对裂缝的非决策屏蔽程度很高：评估标准刚硬，教师被专业规范要求"必须提供有效帮助"，家长的教育参与以过程监控和实时纠错为主要形式。这与当下东亚大型城市中产家庭教育现场的典型特征高度吻合。本文**假设**——而非报告观测到——在这个高屏蔽的环境内存在一小批内部建构空间并未塌缩的学习者。第 4.5 节已写明此格的实证描述仍待田野研究填充；本节交出的判据，正是为了让这个假设可被证实或证伪，而不是为了描述一批已被观察到的人。

本文预测，这些幸存者与他们同龄、同校、同家庭经济背景、同外部教育投入强度的同伴相比，其独特性不在于获得了更多的"有效帮助"——事实上，就外部帮助的量化总量而言，两组并无显著差异。独特性落在另一个此前未被制度性观测的指标上：在他们过去的学习经历中，是否存在规律性出现的、无人干预的"认知独处"时段。

认知独处时段的操作化定义必须从严：它不指物理上的独自一人（一个孩子在房间里独自刷题，但家长的错题报告次日会被推送，这不算独处）。它指时间连续、无目标预设、无外部反馈窗口开放的认知漫游时段——包括且不限于：无任务目的的反复翻阅旧书、未被老师指定作文题的抽屉写作、对无解的难题的独自推敲而不寻求帮助，以及看似与学习无关的长时间涂鸦或实物操作。

本文据此定义两组：**幸存之炉组**＝身处高非决策屏蔽环境、且有规律性认知独处史；**冰封之室组**＝身处同等高屏蔽环境、无此种时段。两组在测试时的任务、指导语、时长、环境完全相同。

5.2 判据一：零启动——本文与全部四家的分岔点

这是本文最要紧的一条判据，也是本文与动机-自我调节一路真正分开的地方。第三节已经指出，"归因于策略"与"归因于能力"这组对照，早在半个世纪前就已被无助型／掌握型研究预测，本文不能把它当作自己的判决性证据。本文的判据必须落在一个只有"装置从未建成"预测得到、而全部动机理论都预测不到的位置上。

这个位置就是**启动本身**。

测试程序（真空启动测试）。撤除所有外部支架——AI 界面、教师即时辅导、结构性任务清单、评价预期。向学习者提供一个隐含丰富问题线索的物理或概念情境，但不给予任何明确的认知任务指令：例如，不告诉他"解题"，只交给他一盒零件、一堆旧报纸，或一个模拟生态系统的无序数据。明确告知：没有正确答案，不计分，不汇报，做什么都可以，也可以什么都不做。然后离开。

主要测量=自发目标表述的计数（言语通道与行为通道并列）。"自发目标表述"的编码准入标准从严：必须是学习者在无任何外部提示下，自行说出或写下的一个指向该情境的、可被判定为"我打算弄清楚／做出／试试某件事"的表述。它不需要正确，不需要可行，不需要完整；它只需要存在。"这堆零件是干什么的"算一条；"我想把它们按大小排排看"算一条；"这个数据里是不是有个规律"算一条。而"老师要我做什么"不算，"这个做完了给谁看"不算——那是在寻找外部指令，不是在生成内部目标。

只记言语会漏掉一整类被试。儿童的内部意图与其言语化之间落差很大，出声思考本身又有反应性；一个孩子沉默地把零件按大小排开，这是一次有目标的操作，按纯言语口径却会被记为零——而这恰好落在本文判据最要命的地方。因此编码须设**行为通道**与言语通道并列：任何可被两名独立编码者判为"指向该情境的、非随机的、有内部组织意图的操作序列"，同样计为一条。判定从严：反复摆弄、拆装、无组织的翻动不计；出现可辨认的分类、排序、对比、复现或试探性改动才计。两通道分别计数、合并判定，并报告编码者一致性。

本文的预测（须以正确的统计形式陈述）：本文主张的不是"冰封之室组人人为零"，而是**该组中存在一个产生零的亚群**——两组的自发目标表述计数因而不服从同一个连续分布：幸存之炉组呈单峰计数分布，冰封之室组在零点出现**超出该分布所能解释的堆积**，应以零膨胀混合模型而非均值比较来检验。签名是零点堆积，主张是亚群存在；若只是整体左移而无零点超额，则本文的"有／无"区分不成立（见第六节第四条）。就该亚群的成员而言：不是少，是零。他们不会犯错，因为他们根本没有进入解决问题的状态；他们展现的不是笨拙的探索或缓慢的进度，而是长时间的被动等待、注意力涣散，或反复寻求任何可被简单执行的明确指令。而幸存之炉组即便在同样陌生的情境中，也会产生若干条——可能很幼稚，可能立刻被自己推翻，但它们存在。

四家最近邻在这一格的预测：

与内隐智力／成就目标理论的对照。这一路预测的分化发生在**失败之后**：实体论者归因于能力并冻结，增长论者归因于策略并更换路径。但这条预测有一个不可省略的前提——**被试得先做了什么，才有失败可归因**。在真空启动测试中，实体论者的典型行为应当是：迅速生成一个过于保守、过早收窄的目标表述（"我就把它们按颜色分一分吧"），做完，然后停下，并在被问及时表示"我不太会做这种没标准的东西"。他有**产物，产物是保守的**。本文预测的冰封之室组**没有产物**。这两种结果在数据上完全可分：一个是低分，一个是空白。若冰封之室组普遍产生保守但存在的目标表述，则本文的"零"不成立，本文应并入这一路作为其极端一端。

与自我决定论的对照。这一路预测，自主性需要长期受挫者在获得完全自主的情境时，内在动机会回升——这正是自主性支持干预有效的机制。真空启动测试提供的恰恰是一个自主性满格的情境：无控制、无评价、无截止。因此自我决定论预测：两组的差异在此情境中应当**缩小甚至消失**，因为压制被解除了。本文预测相反：**差异在此情境中最大**，因为压制的解除不能凭空建起一个从未存在过的起点。这是一条方向相反的预测，且只需一次测试即可分开。

与自我调节学习的对照。这一路预测，元认知策略的缺失可通过教学补足；因此在给予一次简短的策略教学（"你可以先问自己：我看到了什么？我想知道什么？"）之后，两组差异应显著缩小。本文预测：策略教学能提高幸存之炉组的产出**质量**，但无法把冰封之室组从零抬到一——因为被教的是"如何组织一个已经存在的意图"，而缺的是意图本身。可裁决处=策略教学后，冰封之室组的自发目标表述数是否离开零。

与认知卸载论的对照。认知卸载论预测，两组差异应与他们对 AI 工具或外部帮助的依赖程度成正比，因而经过一段短暂的脱敏期后，差异应趋于消失。本文的预测与之相反：差异不会因为工具撤除而消失，因为其中一组学习者的内部结构中从未形成那个发起认知行动的起点本身。脱敏可以消除依赖，但无法凭空建起一个从未存在过的炉膛。

5.3 判据二：归因发生在第几步

第二条判据是第一条的时序补充，用于处理那些确实产生了目标表述、但很快停下的中间态个案。

本文预测，在遭遇自感失败的时刻，两组的报告落在**行为序列的不同位置**上：幸存之炉组的报告落在"我目前的方法或框架不对"（第四步之后），并尝试更换路径；无助型归因者的报告落在"我本身不具备解决这类问题的能力"（同样在第四步之后），并出现快速的行为冻结；而炉膛缺失者的报告落在**更早一步**——"我不知道要开始什么""我不知道你想让我做什么"（第二步之前）。

这三种报告在自陈量表上很容易被混为一谈，因为它们都被受访者说成"我不会"。分辨它们的唯一办法是把报告**锚定到行为时间轴上**：在被试说出"我不会"之前，他是否有过任何一次可观测的、指向该情境的自发操作？有，则属前两类；没有，则属第三类。这条判据的价值在于它把一个内部状态的区分，转化为一个纯粹外部可见的时序问题。

5.4 判据三：训练天花板的不对称

第三条判据处理可逆性。本文预测：给予同等强度、同等时长的自主学习训练后，两组的改善曲线形状不同——幸存之炉组呈现持续上升；冰封之室组在初期有一段由"学会了这个任务的做法"带来的上升，随后进入一个与其早期高帮助密度暴露时长相关的平台期，且该平台期在撤除训练情境、更换任务领域后出现回落。

这条判据需要与本站先导研究的可裁决预测划清界限：先导研究比较的是"建构窗口期内即卸载"与"能力建成后才卸载"两组，其第二轴是**时间**；本文比较的是"高屏蔽×有独处"与"高屏蔽×无独处"两组，其第二轴是**制度屏蔽下的缝隙有无**。两条预测可以同时成立，也可以只有一条成立；它们各自独立可证伪。

5.5 可观测代理的确立，与一处必须撤回的过度主张

本文提出，通过回溯性编纂童年及青少年时期的"认知独处时段"频次与类型，可以独立于学习者当前的能力测试数据之外，形成一个可编码的外部变量。编码框架至少需包含：时段的时间长度、连续性、是否有事后汇报义务、是否有成人到场且处于评估姿态、活动是否由学习者自选。

此处必须明写一件本文初稿做过头的事。初稿曾主张，跨国比较已为此提供了初步的可检验场域，并称"这是一项现在就能跑的检验"。这一主张不成立，本文予以撤回。**现有的国际学业评估数据中并不存在"认知独处时段"这一变量**；它们测量的是课外学习时长、家庭作业时间、课外活动参与，这些量与本文所定义的认知独处在概念上并不重合——一个孩子可以有大量课外活动时间而零认知独处（每一项都有成人到场、有目标、有事后评价），也可以在很短的时间里获得高质量的认知独处。用现有数据跑这项检验，得到的将是一个关于"闲暇时长"的结论，而本文的第二轴恰恰主张闲暇时长不是关键变量。

诚实的表述是：本文的可观测代理需要一份尚不存在的量表。它需要一套回溯性访谈编码方案（用于成人被试）与一套前瞻性日志方案（用于在学儿童），二者都需先做信效度检验。本文把这项工作明写为自己的欠账，而不是把它藏在一句“现在就能跑”里。

另需明写一件与上一段同类的事，以免本文在撤回一项过度主张的同一页上又造出一项。第三节指出，既有研究中被记为零分或无效的被试，在本文框架下正是核心样本——这一判断成立；但“这些数据已经存在、只需重新调取”不成立：**那些研究通常不采集被试的认知独处史**，被试多已匿名，成长环境信息无法回溯补齐。正确的表述是：它指认了一个此前被系统性丢弃的观测类别，而非一份现成可用的数据集。

5.6 判据表格（逐格分条）

判别格：高非决策屏蔽环境（制度对裂缝的系统性屏蔽强）× 学习者面对无模板复杂任务的表现。

第二象限（幸存之炉·高建构空间）：本文预测，此类学习者在真空启动测试中产生多条自发目标表述，失败时的报告落在方法层并尝试更换路径，对无评估模块的任务仍能保持自主投入，训练曲线持续上升。内隐智力理论预测他们应为增长论持有者，无法解释为何持有增长论者在同一制度环境中比例如此之低。自我决定论预测他们的差异应随自主性满格而缩小。认知卸载论预测他们与冰封之室组的差异应随工具脱敏而弥合。技术垄断论无法从文化整体叙事中单独解释其幸存机制。

第四象限（冰封之室·低建构空间）：本文预测，此类学习者在真空启动测试中产生零条自发目标表述，其“我不会”的报告在时间轴上出现于任何自发操作之前，策略教学后仍不离开零，训练曲线呈平台并在换域后回落。内隐智力理论预测他们应产生保守但存在的目标表述并在失败后归因于能力。自我决定论预测他们在自主性满格情境中动机回升、差异缩小。自我调节学习预测策略教学可使其离开零。认知卸载论认为撤去 AI 依赖后其表现应随时间恢复。功绩主体论可将他们解释为规训得不够深或相反——已被完全规训但意义感丧失殆尽，但无法解释为何同一环境中存在两种对立的结果。

这两格的预测在方向上彼此排斥，且全部落在可计数的行为量上。这是本文交给任何愿意接手实证工作的研究者的判决场地。

六、证伪条件

本节不是为论证增加一道保险，而是交出四把可以打开并摧毁本文核心命题的钥匙。一把锁的结构越是精密，它就越应该在自己的设计图纸上标出锁匠自己开不了的那几扇门。本文的承重命题——认知困境的本质不是认知卸载、自我规训或动机受挫，而是内部生成结构（炉膛）在一个高非决策屏蔽的系统中，从未被允许完整建造——将在下列至少一种情形出现时被证伪。

第一条证伪路径：工具脱敏后的结构恢复

本文对认知卸载论的核心反驳，是它无法解释"结构从未建成"与"结构暂时依赖"之间的质的差别。因此，本文的第一个证伪条款是：如果对一批被诊断为"第四象限（冰封之室）"的学习者，进行为期一学期或更长的系统性外部辅助工具戒断（包括且不限于 AI 界面、辅导班模板化解题流程、家长即时纠错），并在此过程中仅提供基本的无框架材料刺激（如题本、书籍、实物操作材料，但不给予任何流程化指导），在学期末，他们在面对全开放问题时的自发目标表述数量、路径更换频率及无评价任务中的自主投入时间这三个指标，与对照组（已被判定为"幸存之炉"类型）之间不再有显著差异，或差异的效应量缩小至一个可忽略的阈值之下——那么，本文关于"永久性绕过导致结构从未建成"的核心判断就是错的。这并非说困境不存在，而是说困境的实质确如认知卸载论所言，是一种暂时的、可逆的功能性依赖，而非发育性的缺失。一个基于现有国家数据的准自然实验已经存在：比较不同国家对数字产品进课堂的禁令实施前后，学生在上游开放性项目中的表现变化，可以提供第一轮检验。

第二条证伪路径：缝隙机制的失效

本文引入第二轴"非决策权力"，将幸存解释为一个发生在制度缝隙中的事件，并把幸存者锚定在"认知独处时段"这个可观测变量上。如果一项严格的追踪研究发现，在排除了社会经济地位、家长教育方式（除开独处容忍度外）、学校质量和学生基线认知水平的混杂后，"认知独处时段"的有无，无法对内部建构空间（通过第五节三项指标测量）产生任何显著的、独立于一般自由活动时间的预测力——即，是"任何形式的休闲活动"而非"特定类型的认知独处"在起作用——那么，本文的第二轴就是一根多余的概念拐杖，本文对幸存机制的解释并未超越常规的教育心理学变量（如闲暇时间、压力水平）。

第三条证伪路径：反向因果的确认

本文所描述的全部困境，有可能既不是外部帮助的形式导致的，也不是非决策屏蔽导致的，而是一个更古老的、与 AI 和技术制度无关的个体差异问题在新时代的重新显现。如果一项对"认知独处时段"的前瞻性研究（即追踪初期尚无此习惯的儿童，在随机分配至不同的教学干预组后，其自发形成认知独处习惯的比例）发现，认知独处习惯的形成，几乎完全被儿童早期的气质类型（如努力控制、感觉寻求）和基本认知风格（如场独立-场依存）所预测，而外部的非决策屏蔽制度环境（高屏蔽组对低屏蔽组）的增量解释力可忽略不计，那么本文的诊断就犯了归因谬误：它不过是给一个古老的个体差异谱系，穿上了一件当代技术批判的外衣。那时，本文应当被撤销，并被重述为对一种特定认知气质在当代教育制度下更易凋零的现象学描述，而非对一种制度性认知病理的发生学诊断。

第四条证伪路径：零启动判据本身的失效

本文把全部判决性重量压在一个存在论区分上：有产物而产物错，与没有产物。这个区分若不成立，本文的一切就坍缩回一个速度谱系。因此本文必须为这条判据自己交出阈值与失败条件。

如果在真空启动测试中，被判为冰封之室的学习者在**足够长的**独处时间之后普遍产生了自发目标表述——那么"零启动"就不是零，只是慢。本文为此预先给出一个操作阈值：**在 60 分钟无指令、无评价、无干扰的独处之后仍无任何可编码的自发目标表述，才算零启动**；10 分钟的短窗口只用于测量启动速度，不用于判定装置的有无。

这个 60 分钟是一个待标定的试点值，不是已有依据的标准——本文没有推导它，也没有试点数据支持它。它需要由一条启动时间的生存曲线来确定：在常模样本上测量首次自发目标表述的等待时间分布，取其高分位（如 95%）为阈值。在标定完成之前，任何据 60 分钟作出的"零启动"判定都只具探索性质。一个判据若把自己阈值的来源藏起来，它的证伪条件就是假的。

若在 60 分钟窗口下，冰封之室组的自发目标表述数分布不再呈现零点堆积，而是与幸存之炉组构成同一个连续分布的低端，那么本文的"有／无"区分是错的，本文应当并入自我调节学习那一路，作为其极端一端的描述，而"炉膛"这个词应当被撤下，因为它承诺了一个不存在的断点。

这条证伪路径比前三条更贴身，因为它可以在一间教室里、用四个小时跑完。

就本文已核验的范围而言，上述四条证伪路径所要求的决定性证据，在已有公开文献中尚未出现。然而，这并不意味着它们在未来不会出现。本文在此处自动注销"本文是唯一可能的解释"这一说法，并将上述四把钥匙交到任何愿意接手此项实证工作的研究者手中。

七、结语：在非决策的遗忘之处

本文的母题，始于一个被长期误诊的教育困境。当批评声浪反复指向 AI 如何削弱记忆力、瓦解文化叙事或催生功绩主体的自我剥削时，它们在反复回应同一个问题：教育的既有好处，被什么弄坏了？在这个追问下，分析的对象永远是那片正在被侵蚀的已知大陆。

本文在此大陆边缘，标绘了一条此前未被命名的海岸线。在那些最令人困惑的案例里——学习者并未展现出任何可被明确归类的病理，却失去了一个无法被任何既有术语捕捉的东西——那片大陆不是被侵蚀了，而是从未被建造。本文称这个缺失物为“认知炉膛”，并论证它的缺席并非起因于 AI 的入侵，而是起因于一个更古老的、遍布于家庭、学校与教育技术规范中的本能：把学习者在裂缝中的独自挣扎，当作一个需要被治愈的伤口来填平。这个本能本身，就是一个长期的、未被追问的非决策。教育史上从未有一份教师手册把“必须保留学生的不解”写入行为规范；这不是一项被否决的议程，而是一项从未生成过的议程。

而在追问这个非决策为何如此顽固时，本文得到了一件它出发时没有预料到的东西。政治学告诉我们，一项议题被挡在议程之外，通常是因为有人从它不被讨论中获益。可是在这里，我们找不到受益方：教师不获益，家长不获益，学生更不获益。挡住这项议题的不是任何一方的利益，而是三套彼此独立、各自都无可指摘的善意规范——教师必须有效帮助，家长必须尽责陪伴，产品必须更快更准——它们在同一个方向上叠加，把同一件事挡在了门外。**没有受益方的偏见动员**，比有受益方的更难松动：因为你没有人可以揭露，也没有人需要被说服。所有人都是对的，而那件事仍然不在议程上。

AI 的意义在此被重新校准。它的危险性不在于它太强大，而在于它恰好与那个古老的填平本能形成了一种史无前例的完美联盟。它把“秒级回应”变成陪伴，“精准诊断”变成理解，“永不疲倦”变成关怀。认知卸载论、功绩主体论、技术垄断论、动机-自我调节理论各自从不同侧面穿透了这间冰封之室，它们的诊断没有错。但它们漏掉了旁边那间幸存之炉，而那些把炉膛保留至今的幸存者，向它们展示了这个判断的边界。辨识这条边界，不是要否定任何人的犀利，而是要让诊断从一把只能打开一扇门的钥匙，变成一张能分辨出四扇不同锁的地图。

本文最后的追问不应在空中盘旋，而应落回我们每个人脚下最具体的现场。在教育学论文里许下宏愿是容易的。真正的风险在另一个地方：当一名教师在明天的课堂上，面对一个卡住了、用求助的眼光看向她的学生时，她能不能咽下喉头那句“简单，老师告诉你”，把沉默多留三分钟？当一个父亲在深夜批改完孩子的作业，发现孩子仍不死心地试一种新的、明显更笨的解法时，他能不能不看表，不问“有什么用”，只是把那条笨路旁边的一块杂物搬开，让它稍微好走一点？

这场认知的保卫战，不发生在教育部的文件里，不发生在科技伦理的会议上，它只发生在这三分钟里，发生在书桌边的那个没有标准答案的深夜。而按本文第四节的分析，这三分钟之所以难，恰恰不是因为教师或家长不愿意——是因为在现行的正当性结构里，这三分钟没有一个可以被说出口的名字。她若沉默，无人会记录她保护了什么；她若开口，所有人都会记录她提供了帮助。**制度能不能为这三分钟留出一块不被计入绩效考核的空地——一块允许被说出口、也允许被写下来的空地**，是本文唯一没有、也无法给出答案的问题。因为它不是一项待办事项——它是一个选择。

附论·与站内先行研究的划界

本文与本站三篇已发表的先行研究共享概念材料，若不划界，读者有理由怀疑本文只是它们的一次重排。四条界线如下。

与《当炉膛从未被建造》（教发生四论之四）。那一篇是本文的直接前身，它命名了"意义引力"与"认知秩序的坍缩"，并论证了"从未被建成"这一判断在概念上不可被既有教育学术语替换。**分工线在阶数**：那一篇交出的是一阶判断——这件东西存在且此前无名；本文交出的是二阶辨别——同一件东西被放进一个由第二轴（非决策屏蔽）张开的二维格中，从而使"炉膛不在"这个全称诊断被切成四种彼此可分的地形。**可裁决处**=第二象限（幸存之炉）的存在与否：先导研究的框架无法预测它，若田野中确实存在一批身处高屏蔽环境却保有内部生成装置的学习者，则第二轴是必要的；若不存在，则第二轴多余，本文应退回先导研究的一阶结论。

与《正当的排斥》（教发生四论之三）。那一篇的分析单位是评价系统，判断是：当"正确"成为唯一合法性尺度，排斥生成是系统维护自身正当性的必要操作。**分工线在层位**：它处理的是排斥为何必然，本文处理的是排斥在制度均质的前提下为何仍有个体差异。第四节已明写，评价刚性被本文排除在第二轴之外，正因为它是 Z 的制度成因而非与 Z 正交的维度。**可裁决处**=在评价刚性完全相同的两所学校间，若个体差异消失，则本文的第二轴不成立，《正当的排斥》的单轴解释已足够。

与《秩序的肌理》（教发生四论之二）。那一篇处理的是既有认知过程被一套外来语法重新组织，即改造。本文第 4.6 节的"被损伤（结构变形）"正对应它。**分工线**=改造与未建成之分：《秩序的肌理》描述的是中等 Z 值带功能性扭曲的状态，本文的核心命题描述的是极低 Z 值的空无。二者在本文辨别格中占据不同位置，可同时成立。

与《认知接触替代与控制补偿锁定》。本文第一节所依据的"认知接触替代"机制来自该文。**分工线在场域与因变量**：该文处理的是 AI 介入家庭学习后亲子之间何种接触被保留，因变量是关系与控制行为；本文处理的是同一替代机制在学习者内部造成的结构后果，因变量是自发目标表述。本文引用它作为机制前提，不重复其论证。

附论·本文与初稿的差异（编辑说明）

本文由投稿原稿经编辑性提升而成。为使读者能够独立判断哪些部分归于原作者、哪些归于编辑增补，此处逐项列明。

原稿已有：第一至二节全部，第三节 3.1，第四节 4.1–4.3、4.5–4.6，第五节的判别格设定与两组定义，第六节前三条证伪路径，第七节结语的前四段。

编辑增补：① 第三节 3.2（动机-自我调节一路的消解、“缺席是否可逆”与“归因发生在第几步”两条分离线）与 3.3（与生成性失败／合意难度／帮助两难的分工）；② 第四节 4.4（正当性型非决策，及其对卢克斯／布迪厄／规范性同构的地盘让出）；③ 第五节 5.2 的零启动判据、零膨胀混合分布口径、言语与行为双编码通道、四家逐一对照，5.3 判据二，5.5 中对跨国数据与旧数据回溯两项主张的撤回；④ 第六节第四条证伪路径、60 分钟阈值及其待标定声明；⑤ 结语中关于“没有受益方的偏见动员”与“这三分钟没有名字”两段；⑥ 本页两则附论。

增补部分的自复核（2026-08-03，第二轮）：增补完成后做过一次自复核，改掉了增补自身的三处硬错与两处不一致——（甲）原写“没有一支处理装置是否存在这个问题”，不成立，自我决定论的“无动机”正是该范畴，分歧已改锚在可逆性上；（乙）“正当性型非决策”原被写成对非决策概念的修正，实际上卢克斯、布迪厄与规范性同构已占其大部，已收窄为“多规范叠加”与“补救不对称”两处残余；（丙）原写零产出白卷“只需重新调取旧数据”，与本文自己撤回的跨国数据主张同病，已一并撤回；（丁）5.2 的存在论表述与第六节第四条的分布语言口径不一，已统一为零膨胀混合模型；（戊）5.1 原以观测语气陈述幸存者，与 4.5 承认此格尚无实证描述相冲突，已降为假设。此外零启动的编码原只走言语通道，会漏掉沉默操作者，已补行为通道。**本条列出这些，是因为一篇把证伪条件当作主要产出的文章，没有理由把自己的订正史藏起来。**

原稿删改：删去两处未定义即使用的内部符号标记；将摘要中“第二轴来自社会学与组织理论”改为“来自政治学”（巴赫拉克与巴拉茨 1962 年的论文发表于《美国政治学评论》）；补完被截断的结语与参考文献。

按本站规程，改稿者不得为自己改过的文本盖认证分，故页首所标创新智商仅对应投稿原稿，上述增补部分不计入。

参考文献

1. Bachrach, P., & Baratz, M. S. (1962). Two faces of power. *American Political Science Review*, 56(4), 947–952.
2. Bourdieu, P. (1977). *Outline of a Theory of Practice* (R. Nice, Trans.). Cambridge University Press.
3. DiMaggio, P. J., & Powell, W. W. (1983). The iron cage revisited: Institutional isomorphism and collective rationality in organizational fields. *American Sociological Review*, 48(2), 147–160.
4. Lukes, S. (1974). *Power: A Radical View*. Macmillan.
5. Meyer, J. W., & Rowan, B. (1977). Institutionalized organizations: Formal structure as myth and ceremony. *American Journal of Sociology*, 83(2), 340–363.
6. Bjork, E. L., & Bjork, R. A. (2011). Making things hard on yourself, but in a good way: Creating desirable difficulties to enhance learning. In M. A. Gernsbacher et al. (Eds.), *Psychology and the Real World* (pp. 56–64). Worth Publishers.
7. Deci, E. L., & Ryan, R. M. (2000). The "what" and "why" of goal pursuits: Human needs and the self-determination of behavior. *Psychological Inquiry*, 11(4), 227–268.
8. Dewey, J. (1910). *How We Think*. D. C. Heath & Co.
9. Diener, C. I., & Dweck, C. S. (1978). An analysis of learned helplessness: Continuous changes in performance, strategy, and achievement cognitions following failure. *Journal of Personality and Social Psychology*, 36(5), 451–462.
10. Dweck, C. S. (1986). Motivational processes affecting learning. *American Psychologist*, 41(10), 1040–1048.
11. Illich, I. (1971). *Deschooling Society*. Harper & Row.
12. Kapur, M. (2008). Productive failure. *Cognition and Instruction*, 26(3), 379–424.
13. Koedinger, K. R., & Aleven, V. (2007). Exploring the assistance dilemma in experiments with cognitive tutors. *Educational Psychology Review*, 19(3), 239–264.
14. Piaget, J. (1952). *The Origins of Intelligence in Children*. International Universities Press.
15. Postman, N. (1993). *Technopoly: The Surrender of Culture to Technology*. Vintage Books.
16. Risko, E. F., & Gilbert, S. J. (2016). Cognitive offloading. *Trends in Cognitive Sciences*, 20(9), 676–688.
17. Sparrow, B., Liu, J., & Wegner, D. M. (2011). Google effects on memory: Cognitive consequences of having information at our fingertips. *Science*, 333(6043), 776–778.

18. Vygotsky, L. S. (1978). *Mind in Society: The Development of Higher Psychological Processes*. Harvard University Press.
19. Zimmerman, B. J. (2000). Attaining self-regulation: A social cognitive perspective. In M. Boekaerts, P. R. Pintrich, & M. Zeidner (Eds.), *Handbook of Self-Regulation* (pp. 13–39). Academic Press.
20. 韩炳哲 (2019). 《倦怠社会》 (王一力译) . 中信出版集团.
21. 王德生 (2026). 《认知接触替代与控制补偿锁定：生成式AI介入家庭学习的机制模型与可裁决研究设计》. SDE Universes. (站内·AI时代的教育)
22. 王德生 (2026). 《正当的排斥：为什么教育系统越正当，越在发生学上饿死认知》. SDE Universes. (站内·教发生四论)
23. 王德生 (2026). 《秩序的肌理：当评价结构被内化为认知的骨骼》. SDE Universes. (站内·教发生四论)
24. 王德生 (2026). 《当炉膛从未被建造：为什么「教发生」的改革会撞上沉默之墙》. SDE Universes. (站内·教发生四论)