

构成与可构成者：重勘创新的发生位

我们用一套只能识别“物”的筛子，筛掉了让新的“认识者”得以生成的条件

胡志英 著 · SDE 学员 · 创新理论与科学史 · 2026年7月27日

摘要 本文将“创新”的追问从产出端移向发生端，提出一对不可还原的区分：一边是已被明确表述和编码、构成一个时代公共认知库存的“构成性知识”；另一边则是“可构成者”——不是一个可被命名为“创新方法”的成品，而是一个认识者在特定认知压力下，将自身重新构成为能够看见新问题的人的事件性过程。通过将狭义相对论的诞生置于与洛伦兹电子论的对称比较中，本文从“悬置与边缘失效”“问题漂移”“对话性打断”与“裁决性肯定”四个环节，展示可构成者发生的非必然结构及其失败形态。在理论定位上，本文正面论证“可构成者”不可被库恩的“格式塔转换”所还原——前者不是一次“学会看见新格式塔”的瞬间，而是一场在对话压力下被迫发生的、使认识者丧失旧有“同一性基底”的不可逆断裂；并通过与阿吉里斯“双环学习”和当代认知科学类比推理范式的对勘，进一步廓清概念边界。基于这一生成机制诊断，本文追溯当代创新管理中“三重滤网”的形成史，揭示其如何从合理的纠偏机制耦合为对悬置、漂移与裁决性肯定的系统性排斥，同时指出滤网自身的矛盾正在积聚瓦解自身的张力。最后，本文为辨识与守护可构成者提出一个自反性的防守框架，将“守护”的落脚点从呼吁文明自觉移向暴露守护条件自身的不可持续性，并在文末将可构成者的蜕化条件作为理论自身的证伪命题纳入反驳节，而非安全阀。

关键词：构成性知识；可构成者；认知重构；裁决性肯定；问题漂移；格式塔转换；爱因斯坦；洛伦兹

一、把追问从产出端移到发生端：一个未经检视的预设

我们生活在一个对“创新”既极度渴望又深受其困的时代。庞大的创意产业——从企业创新管理软件到人工智能驱动的想法生成器——正以前所未有的精密度承诺可以加速、管理和量产“金点子”。然而，一个悖论性的事实始终悬挂在这股热潮顶上：在创意工具空前发达、创新投入空前巨大的近三十年里，人类在基础科学范式、重大技术架构方面的突破性创新，并没有呈现同步的爆发式增长。我们拥有了更多的新产品、新应用和新的商业模式，但那种真正改写游戏规则的东西，依然极为罕见。

这一悖论激发了大量反思。其中一条有力的分析路径，是指出当代创新工具和管理体系所擅长的是“组合式新颖”——在既定概念空间中重新排列已知元素，产出形似创新的新组合；而它们所无法触及的，是那种改变问题框架本身的“发生式创新”。这一区分富于启发，但它自身也预设了一个值得追问的前提：它仍然把创新看成是两种并列的“知识产出类型”——一种是组合出来的，一种是发生出来的。本文要推进的，是更深一层的追问：当我们把创新当作“特殊的知识产出”加以管理、评价和催生时，我们是不是已经悄悄地把“知识”和“知道者”当成两个可以分开处理的东西，并且默认那个“知道者”是一个稳定的、随时待命的认知主体——只要知识工具足够好，他就能产出更好的创新？

这正是本文试图打开的新问题域。我将提出一个更精确的区分来取代旧二分法。一边是“构成性知识”：那些已经被明确表述、传授、编纂、输入算法、写进教科书和KPI指标体系的东西。它包括公式、概念、方法步骤、评价标准、案例库和最佳实践。它们构成一个时代、一个组织、一个学科的公共认知库存。

另一边，我命名为“可构成者”。这里必须从一开始就划清界限：本文所称“可构成者”，不是指一种稳定的认知能力，不是某些人天生具备的特殊禀赋，也不是一个可以被拎出来命名为“创新方法”或“思维模型”的成品。它指的是这样一个事件性过程：一个认识者，在长期浸淫于特定的构成性知识体系并与之反复碰撞的过程中，逐渐被逼到了某个认知边缘——在这个边缘处，让他继续认知下去的唯一方式，不是调用更多的构成性知识，而是他自己的认知框架本身被迫发生重组。这个过程不是构成性知识推导出来的，但也绝不是与知识无关的“纯粹灵光”；它恰恰是在与构成性知识最深度的交织中，从那些知识无法覆盖的缝隙里发生的。它不是“什么”——不是可以被收纳进公共库存的认知产出；它是“谁”的问题——是那个认识者在知识中重新构成自身的事件。

这样一来，当代创新的根本困境，就不再仅仅是“我们太强调组合，忽略了发生”。那仍然把创新当成两种并列的知识类型，仍然在追问“产出”而不是追问“产出者如何被构成”。真正的困境更深一层：我们所有的管理体系、评价指标和创意工具，都建立在一个未经检视的预设之上——即“创新是一种特殊的知识产出，认识者是给定的”。因此这套体系天然地只能处理“构成性知识”层面的操作；而真正的突破，恰恰不是一个“新知识”替换“旧知识”的事件，而是一个“可构成者”——一个被新问题逼到认知边缘的认识者——在给定制度的缝隙中，将自己重新构成为一个能看见新东西的人。我们用一套只能识别“物”的筛子，无意间筛掉了让新的“认识者”得以生成的几乎全部条件。这才是悖论的核心。

下文从四个层面展开。第二节确立核心区分，并完成与赖尔、波兰尼的最初划界。第三节进入本文的论证重心：通过爱因斯坦与洛伦兹的对称比较，展示可构成者的发生机制及其失败形态，并在其中正面完成与库恩“格式塔转换”的不可还原性硬论证，同时回应“审美

裁决是否会滑向天才心理学”的质疑。第四节追溯当代创新管理“三重滤网”的发生史。第五节迎战组织学习理论中“双环学习”和认知科学中类比推理范式这两个最强劲的对话者。第六节为辨识与守护提出防守性框架，并对其可能退化为新教条的风险保持自觉。第七节正面回应“伪构成”“AI创新”与“精英主义”等严肃反驳，并将可构成者的蜕化条件作为证伪命题纳入反驳节，而非安全阀。

二、构成性知识与可构成者：一个不可还原的区分

2.1 构成性知识：公共认知库存的特征与边界

构成性知识，是任何一个时代、学科或组织内部已经被明确编码、可被传授、可被写入手册和算法系统的知识。它不限于“事实”或“公式”，还包括方法论的步骤（如何设计一个随机对照试验）、评价标准（什么才算一篇好论文）、成功的先例（那些被奉为经典的案例研究），以及隐含着价值判断的分类体系（如将疾病划分为“器质性”与“功能性”）。其本质特征是已经构成了固定的概念节点、明确的关系连线与相对稳定的评价规则。正因如此，它被高效传播——一个从未见过细胞的医学生可以在一周内掌握细胞病理学的基本概念——也能被算法处理，知识图谱和AI辅助诊断都建立在这一前提上。

在熊彼得（Schumpeter, 1912）的“新组合”理论中，创新被定义为“执行新的组合”，即把已有的生产要素、技术手段和市场渠道以未曾有过的方式重新配置。这是一个深刻的洞见，它准确捕捉了人类创新史中大量真实发生的过程：古腾堡的印刷机、共享经济平台都是这种新组合的产物。熊彼得实质上是将创新视为构成性知识的重新部署：知识是给定的元件，创新者是已经给定的“组合者”，其认知结构本身不需要被追问。

本文要指出的，正是这个框架的边界。熊彼得的理论完美解释了“组合式新颖”，但无法解释另一种情况：当创新涉及的不是“元件如何重新连接”而是“元件的定义本身被改变”时，那个改变是如何在一个具体的认识者身上发生的。因为要重新定义“什么是元件”，认识者必须先经历一个过程——在这个过程中，那些曾经向他保证了“这就是元件的正确分类”的整个知识框架，在他体内发生了断裂和重建。这不是“更好的组合”，而是“组合者自己被重组”。这就是“可构成者”的切入点。

2.2 可构成者：在认知边缘处被逼出的重构

可构成者不是一个知识品种，而是一个过程。一个认识者在长期浸淫于某个构成性知识体系的过程中，逐渐感到这个体系在某处“说不通”——不是因为发现了体系内部的逻辑漏洞，而是因为他面对某个具体情境时，给定知识所提供的那套问题、分类和评价标准，无法覆盖他正在经历的不安。这种不安最初是模糊的、无法言说的，因为它要陈述的“问题”本身就依赖于给定知识的语言，而恰恰是这套语言在那一点上失去了覆盖力。

此时他的同事用给定知识向他解释为什么“这没什么不对劲的”，他明知那些解释在技术上是正确的，却仍然感到不安。这种不安逼着他不断试错——不是试方法，而是试问题本身。每一次试错都暴露出给定知识的某种分类、某个预设或某条评价标准在应对这个特定交织时的不自洽。这个过程是漂移的：他最初以为问题出在A处，投入大量时间试图修补A，然后发现A背后预设的B才是真正的麻烦，进一步追到B所依赖的C。在这个过程中，他不是“学习新知识”，也不是“发现新事实”；他是在将自己从那个给定知识所构成的认识者，重新构成为一个能够看见那些不显现在给定知识内部的问题的人。这就是为什么我称之为“可构成者”，而非“洞见”或“直觉”——后者暗示着某种天赋的心理能力，而我要指出的，是一个在认知压力下被迫发生的认知主体的重构。

这一区分与库恩的“范式转换”共享研究对象——都是那些改变了问题定义和评价标准的根本性认知变革。但我将在第三节中通过爱因斯坦与洛伦兹的对称比较，正面论证：库恩的“格式塔转换”叙事无法解释可构成者过程中的两个核心环节——“问题漂移”和“裁决性肯定”——从而确立“可构成者”不可还原的理论收益。此处先行悬置。

2.3 先行划界：对决赖尔与波兰尼

在进入库恩之前，我先处理两个必须切割的经典概念。

第一个是赖尔的“知道如何”与“知道那个”的区分。赖尔（Ryle, 1949）指出，智能行为不能被还原为对命题性知识的内部遵循——一个会下棋的人并非先在脑中默诵规则再执行动作，他的智性体现在其行为本身之中。从表面看，构成性知识似乎对应于“知道那个”，而可构成者所描述的“在规则失效时发生的认知重构”似乎与“知道如何”的领域重叠。

但赖尔的核心关切是行为的智能性质，其分析单元是行为者展现出的“知道如何”。在这个框架中，认知主体本身是稳定的：一个骑自行车的人无论技能多么精湛，他始终是“那个掌握了骑车技能的人”。他的“知道如何”是一种他已拥有的能力。可构成者指向的则是性质上不同的问题：它关心的不是一个认识者“拥有什么能力”，而是那个认识者“如何变成了一个不同于此前的人”。当爱因斯坦在1905年松开“同时性是绝对的”这一预设时，他不是获得了一项新的知道如何的技能——他是经历了一场认知身份的断裂与重建。那个由

牛顿力学和经典电磁学共同构成的“认识者爱因斯坦”，在那一瞬间被一个不同的认识者所取代。赖尔追问的是同一主体在不同情境中如何展现智能；本文追问的是主体自身如何在认知极限处被重组。“可构成者”不可被还原为“知道如何”的子类。

第二个是波兰尼的“默会知识”。波兰尼（Polanyi, 1958）揭示的是在所有认知活动中都存在着“我们知道的多于我们能说出的”个人系数——如骑车时身体对倾斜度的实时感应或老医师瞥一眼病人就涌起的“不对劲”感。这是一种稳定的、可调用的个人基底，它始终在那里。

可构成者则不同。它不是“知道得更多”，而是“被重新构成”。在断裂发生前，认识者首先是被给定知识构成的——他学会了用那套概念看问题。断裂发生时，他失去的不只是一个错误答案，而是一整套使他能够稳定认知世界的支撑结构，跌入“不知道该怎么看”的深渊。默会知识在他跌落时可能起到缓冲作用——爱因斯坦对物理理论的审美直觉无疑在他做出裁决时发挥了作用——但它本身不解释断裂的发生。波兰尼关心的是知识的个人根基；本文关心的是那个根基在特定极限情境下被连根拔起、重新生长的过程。两个问题共享边界地带，但追问方向根本不同。

三、可构成者的发生与不发生：爱因斯坦与洛伦兹的对称解剖

如果可构成者的区分不是空的，它就必须能解释同一个历史情境中，为什么甲发生了重构而乙没有发生。本节在爱因斯坦狭义相对论的诞生之外，引入亨德里克·洛伦兹——这位比爱因斯坦年长二十六岁的荷兰物理学家，同样深深浸淫于经典电动力学与牛顿力学的矛盾，同样积累了数十年的理论挣扎，同样拥有欧洲学术界的最高声誉和最充分的“悬置条件”，却最终未能经历可构成者的重构——作为对称的失败案例。通过两者的对照，逼出可构成者发生的真正的必要条件——那些无法被“格式塔转换”“默会知识”或任何方法论程序所吸收的环节。

3.1 对称案例的合理性：洛伦兹为何是最恰当的对照者

选择洛伦兹作为对照者，基于以下对称性。第一，面临同一组矛盾：两人都在麦克斯韦电动力学与牛顿力学如何兼容这一核心问题上投入了长期精力。第二，深度理论沉浸：洛伦兹曾在1890年代独立发展出“长度收缩”假说（后称洛伦兹-菲茨杰拉德收缩）来解释迈克耳孙-莫雷实验的零结果，并于1904年发表了一组坐标和时间变换方程——后来被庞加莱命名为“洛伦兹变换”——在数学形式上与爱因斯坦1905年给出的变换完全等价。第三，两人都“看到了同样的数学”：1905年之前，洛伦兹和庞加莱已经触摸到了狭义相对论的几乎全部数学骨架。第四，充分的悬置条件：洛伦兹是莱顿大学的物理学教授，享受欧洲学术界的最高尊重和完全的学术自由，没有发表论文的KPI压力，没有需要取悦的基金评审人。第五，两人都与同行有深度对话：爱因斯坦有贝索，洛伦兹与庞加莱、维恩等一流物理学家保持着长期的通信和讨论。

条件如此对称，结果却截然不同。到1905年，爱因斯坦突破了经典框架，而洛伦兹——尽管自己亲手写出了新变换方程——始终将“长度收缩”和“局域时间”仅仅解释为真实的物理过程在以太参照系中的“表象”，而以太本身和绝对时空仍是不可动摇的实在。一个已经握住了全部数学钥匙的人，为什么就是转不过那道门？

3.2 悬置与边缘失效：两条分岔的路径

爱因斯坦的“悬置”之所以有效，其前提不仅是有时间，更是有位置。他身处学术圈边缘——瑞士伯尔尼专利局三级技术员——这一边缘性使他免于被专业共同体的“什么算正经问题”的嗅觉所规训。那个十六岁追光的疑惑，在资深物理学家看来是一个无意义的假想实验——没有人在速度达到光速时活着看到任何东西，它不值得认真对待。但恰恰因为爱因斯坦尚未完全内化这一专业纪律，他允许那个疑惑在脑中存留了十年。相比之下，洛伦兹在1890年代已经是一个功成名就的资深物理学家，他对麦克斯韦方程组的尊重和对牛顿绝对时空的信念，是同一块地基的两面。当他面对迈克耳孙-莫雷实验中光速不变的反常时，他的构成性知识框架立刻提供了一个“合理”的容纳空间——引入以太、给以太赋予收缩的力学性质。问题是“为什么以太会收缩”，而不是“时间和空间本身是否需要被重新定义”。洛伦兹的悬置是在给定框架内部的悬置，他从未被逼到“这个框架本身可能错了”的那个边缘。

3.3 问题漂移：失败是否必然触及底层预设？

在反复的试错中，爱因斯坦和洛伦兹同样经历了失败。爱因斯坦试图用以太理论和传统力学手段去解释追光悖论，每一个模型都在某个关节崩掉；洛伦兹试图在经典框架内为变换方程赋予一个符合牛顿力学的存在论基础，每一次努力都使以太的性质变得更加古怪——它必须同时是绝对静止的、又是能被物体穿过而不产生可探测阻力的、又能在穿过它的物体内部产生长度收缩。两套失败序列却走向了完全不同的终点。

对洛伦兹而言，失败始终是物理问题。他每一次被逼退一步，只是把模型调整得更精巧，从未让失败去触碰“绝对同时性”和“绝对空间”这些最底层的概念预设。在他1904年的论文中，他给出了几乎完整的变换方程，但仔细看他的解释策略：他坚持存在着一个“真实

的”绝对时间和一个由以太决定的绝对参照系，他的变换只是说，对于运动中的观察者来说，测量出的时间是“局部的”、表观的时间——不是真正的时间，而是时钟被运动影响之后显示出来的读数。

对爱因斯坦而言，失败却是问题重定义者。每一次以太模型的崩掉都不只是“这个模型不行”，而是“这类模型之所以屡次崩掉，是否暗示着所有模型底下的那个预设本身就错了？”十年间，他最初以为问题是如何在数学上调和麦克斯韦和牛顿，到1905年春天，经历了无数次崩掉之后，原问题已不再是问题，一个新问题从那些崩掉的缝隙中浮现——“同时性”这个词本身是什么意思？

这里的区别不是智力的区别。洛伦兹丝毫不比爱因斯坦笨。区别在于洛伦兹的认知框架内部没有一个机制能让失败去触及“同时性是绝对的”这一底层预设，因为这套框架本身就是用“绝对时间”作为定义“测量”“实验”“客观物理量”这些基本概念的土壤。要质疑绝对时间，就等于要质疑他自己作为一名物理学家迄今全部工作的那个地基。而爱因斯坦的边缘性在这一点上恰恰成全了他：他所浸淫的构成性知识还没有在他体内形成一套将“绝对时间”当作“不可触犯的硬核”来保护的免疫系统。于是失败可以一路往下追，追到最底层，追到那扇从未被人敲过的门，然后撞上去。

3.4 对话性打断：贝索的功能与洛伦兹对话的失效

1905年春天，爱因斯坦与贝索的那次散步式的长谈，成为压倒性的一笔。这次谈话之所以是必要的打断，不是因为贝索有更深的物理知识（他作为物理学家成就平平），而是因为谈话打破了爱因斯坦独自在内部酝酿问题的模式。当他不得不将自己的混沌不安翻译成话语、抛给一个坐在面前的特定听众时，他被逼着从一个“可能被反问”的位置重新审视那个交织了他十年的问题。就在这种被迫说清的压力下，一个此前从未被他自己意识到的预设——“发生在不同地点的两个事件是否同时，取决于观察者的运动状态”——第一次被他自己的耳朵听见了。它不是被“发现”的，而是被那个说清的行为逼出来的。

对比洛伦兹的对话史。洛伦兹与庞加莱等人有过大量深入讨论，为什么这些对话没有产生类似的打断效果？因为洛伦兹的对话对象和他共享同一套底层预设。当一个洛伦兹和一个庞加莱坐在一起讨论“变换方程的存在论基础”，双方都默认“绝对时间是存在的，变换只是在调整表观测量”。这种对话是高度专业的、极其深入的，但它不会制造打断——因为对话的双方被同一套构成性知识构成，他们互相确认了对方的预设，而不可能用一个外部的追问去撼动那个预设。

可构成者讲的发生，不是在孤独的顿悟中完成的，也不是在任何对话中都能完成的。它是在一个认识者试图向另一个认识者说清自己的混沌时，被那个说清的过程本身逼出的断裂。这个他者的功能不是提供答案，不是施加评判；他的功能是仅仅作为一个“另一个人”在场，迫使认知者将原本在自己内部混沌流动的不安压成一个可以被听见的形式——而就在这个压紧的瞬间，一个从未被质疑的预设，在说话者自己的耳朵里，第一次被听见了。

3.5 裁决性肯定：一个认识论机制，而非心理学叙事

接下来要解剖的是这一节最关键的一笔，也是回应审稿人关于“如何避免滑向天才传记学”质疑的核心所在。在对话性打断中，爱因斯坦意识到所有的麻烦都来自同一个预设——“同时性是绝对的”。但意识到这一点，不等于做出裁决。这里的认识论困境是：没有任何已知实验可以判决性地告诉他“光速在一切惯性系中不变”比“存在绝对同时的以太系”更正确。两套体系各自内部严密自治，各自被大量观察支持，在仅有的反常面前——迈克尔孙-莫雷实验——洛伦兹模型同样可以给出一种解释（长度收缩），不需要放弃绝对时间。从这个意义上看，并不存在一个“外部证据”逼着爱因斯坦选择光速不变、放弃绝对同时性。两条路都通。他必须做出裁决，但他找不到让既有科学共同体承认的“理由”来做出这个取舍。

这时，他的审美信念——自然界应当表现出内在的统一性、物理定律应当对所有惯性的观察者平等适用——从模糊的偏好上升为裁决性原则。这里的危险是：如果裁决最终依赖的是个人的、内部的审美信念，那么可构成者理论是否只是在用更精致的语言复述一种关于“天才个体如何凭直觉做选择”的心理学叙事？

要回应这个质疑，必须推进到认识论层面。裁决性肯定不是一次任意的主观偏好的调用，而是长期认知压力通过审美形式结出的必然结晶。这里需要引入一个更精确的刻画。

一个认识者浸淫于构成性知识多年，经历了无数次失败试错后，其认知结构并不是一张“把所有选项平铺在面前”的白纸。长期的失败已经在某些选项上堆积了沉重的否决记录——洛伦兹的那条以太道路，爱因斯坦试了不止一回，每一回都崩掉。这些否决记录不会直接替他选择光速不变，但它们使得整个问题的可行解空间被压出了一个特定的“压力地形”：有的方向是上坡路，每走一步都必须克服越来越不自然的假设（以太必须同时具备互不相容的性质）；有的方向是下坡路，一旦松开一个特定的预设，原先的堆积在以太上的那些纠结就全部消解。审美不是在这个地形之外的另一个东西——审美恰恰是认知者在长期浸淫中形成的对整个地形的体感。当爱因斯坦说他“相信自然应当是简单的、统一的”，他不是在表达一种纯粹的个人品位；他是在运用一种被长期训练出的、对“这条路上将有更多崩塌还是更少崩塌”的整体探测感，在尚无法用构成性知识将其论证出来之前，先做出了选择。

但这个刻画引出一个更深的问题：这种“压力地形”难道不是爱因斯坦个人的独特心理学产物吗？洛伦兹难道没有他自己的“压力地

形”——对他来说，“放弃绝对时间”恰恰是最陡的上坡路，因为它会动摇他作为物理学家全部工作所依赖的存在论地基？没错。这正是裁决性肯定的第二个关键特征：裁决之所以是裁决，恰恰因为存在两个不可比较的“压力地形”——来自不同的认知构成史，在不同的方向上堆积了否决记录和美学承诺。爱因斯坦的边缘史让他在以太道路上积累了压倒性的否决，洛伦兹的资深史让他在“放弃绝对时间”这一选项上堆积了无法承受的压力。两人的压力地形并不对称。裁决不是从外部仲裁两种地形的优劣，而是每一个认知者在自己的地形中抵达了唯一的稳定解。

一旦澄清这一点，裁决性肯定就从“天才传记学中的审美瞬间”被重新定位为“特定认知构成史在边界条件下的必然出口”。它既依赖于一个认识者长期浸淫于构成性知识所形成的内在压力结构，也依赖于这个压力结构在经历了对话性打断后被迫处理某个此前未被质疑的预设时，所触发的不可逆选择——在那次选择中，旧的认识者被留在了悬崖的另一侧，新的认识者开始用一套从未有过的眼光看东西。裁决是一个断裂性的、不可逆的自我选择——它重新定义了“谁”在看。

3.6 与库恩“格式塔转换”的不可还原性硬论证

有了爱因斯坦与洛伦兹的对称比较，现在可以正面处理本文最核心的理论对话：可构成者与库恩“格式塔转换”的关系。库恩（Kuhn，1962）在描述科学危机中的个体认知者时，引入了“格式塔转换”这一比喻：科学家在危机中突然“看”到了一张全新的图景——就像面对同一张鸭兔图，上一秒看见的还是鸭子，下一秒突然只能看见兔子。库恩还说，在转换前后，科学家的“世界”改变了——他们生活在不同的世界里。

这个描述与本文所说的“认识者自身的重构”高度相似。本文是否只是在用一套更富戏剧性的语言重述库恩？我将通过三个不可还原的差异来论证不是。

第一，格式塔转换是一次性的“整体切换”，而可构成者是一个经历过“无法看清”的深渊的过程。库恩的格式塔转换，其核心现象学特征是“瞬间整体切换”——科学家没有经过一个“既不是鸭子也不是兔子”的中间状态，他的感知从格式塔A直接跳到了格式塔B。但爱因斯坦的案例揭示出完全不同的时间结构：从十六岁的模糊疑惑到1905年春天的裁决性肯定，中间是长达十年的悬置、反复试错、拖延、回避、被另一个问题打散又回来——而且其中相当长的一段时间，他根本不知道自己要往哪个方向走。他不是在A和B两个已经稳定存在的选项之间切换；他是在一片“不知道该怎么看”的黑暗里，靠一次又一次的失败崩掉，慢慢朝向一个尚未成形的B'摸索。在这十年的大部分时间里，他既没有“牛顿的绝对时空”这块完整的格式塔可看（因为那个格式塔时不时被追光悖论戳破），也没有“相对论时空”这另一块完整的格式塔可看（因为后一个七年之后才被他自己造出来）。库恩的“格式塔转换”无法容纳这个漫长的、漂移的“中间态”——在这个中间态里，认识者不仅不确定答案，甚至不确定问题。

第二，“格式塔转换”保留了一个同一的感知主体——前后是“同一个人”看见了不同的东西；而可构成者的重构恰恰动摇了那个使“同一个人”得以成立的根基。在库恩的叙述中，科学家在危机中经历格式塔转换，从旧范式下看到了一个新世界，但他仍然是“那个科学家”——同一个职业身份、同一个物理学家共同体的成员，只不过他的信念被重新分配了。这正是洛伦兹没有发生重构的结构性原因：他太完整地专业共同体构成了，那个“同一个物理学家”的身份是他认知世界的地基，他无法放弃它与放弃绝对时间同时发生。而对爱因斯坦来说，1905年春天的那次裁决，使得“在此之前由牛顿力学与麦克斯韦电动力学共同构成的、能够在专业共同体内被识认为‘物理学家爱因斯坦’的那个认识者”，不再能够继续存在。这不是一个人看见了两张不同的图；这是一个人在裁决发生的那一瞬，被切成了两个不连续的认识者——旧的那个留在裁决的这一侧，丢掉了自己认知世界的一块基石。可构成者讲的不是知觉层面的整体转换，而是存在论层面的认知身份的不可逆断裂。

第三，也是决定性的：库恩的“格式塔转换”无法解释“对话性打断”为什么是必需的环节。诚然，库恩重视科学共同体和同行交流，认为它们是常规科学运作的基础。但在他的危机叙事中，格式塔转换本身归根到底是个体内部的心理事件——是一个科学家面对反常积累到临界点后发生的“看法的根本改变”。他者在这个机制中没有结构性角色。可构成者的论证则指出：没有贝索的那次打断——不是作为提供信息的对话者，而是作为迫使爱因斯坦将混沌压成话语的“另一个人”——那个被十年悬置和无数次失败所积压的张力，很可能不足以自行结算为一次不可逆的裁决。回顾洛伦兹：他同样经历了深入的专业对话，但那些对话的对方与他共享同一套底层预设，因此从未形成打断。对话性打断要起效，需要一个特定的他者位置——一个能够让他者的“可能听不懂或可能反驳”的压力重新配置说话者自身认知压力的位置。洛伦兹的对话史中缺少的，不是对话的数量或深度，而是这个特定的结构性位置。这一环节是库恩的理论工具箱中完全没有的，也是任何将“创新”归结为个体内部认知过程的框架所无法覆盖的。

基于这三重差异，“可构成者”与“格式塔转换”的不可还原性得以确立：可构成者不是格式塔转换的“个体版重述”，也不是在库恩的基础上“补充了一点微观过程”。它描述的是一个性质上不同的发生结构——认识者的重构不是在看见新格式塔时完成的，而是在失去旧格式塔被打断中、在对一个他者说清自己的混沌时，被那个说清的过程本身逼出的不可逆断裂。因此，“可构成者”的理论收益在于：它将原先被库恩归到“个体心理事件”黑箱中的、发生在断裂瞬间前后的那个过程，打开为一个可被严整分析的、涉及认知者身份变迁的发生结构；而这个结构在组织学习和认知科学的主流范式中，至今没有一个等价的概念。

四、三重滤网：排斥逻辑的发生史而非病灶诊断

本节追溯当代创新管理中“三重滤网”的发生史。所谓三重滤网，是分别安装在创新活动的“入口—过程—出口”三个关节处的一组筛选机制：入口处，申请书、项目评审和KPI规划要求创新必须提前以构成性知识的形式陈述为“一个可解决的问题”；过程中，节点审查、中期汇报和阶段性成果考核要求创新必须在给定时间窗口内展示可被构成性知识识别和量化的“进展”；出口处，市场检验、引用率、转化效益等指标要求创新必须在给定评价标准下被判定为“成功”。这三重滤网本身不是由一个中心化权力自上而下设计出来的阴谋。它们各有自己的合理起源，因此批判它们不能简单地将其诊断为“排斥逻辑”——那会把一个发生过程冻结为一个诊断标签，遮蔽对这些滤网发生史的理解和对其内在矛盾的把握。

4.1 从合理纠偏到滤网耦合：三重滤网的发生史

入口滤网（问题必须可陈述）的合理起源在于对“空想”的纠偏。在二十世纪中叶以前，大量科研经费投入了目标模糊、论证空洞的“宏伟设想”。从二战后的美国海军研究办公室到1950年代的国家科学基金会，评审制度逐步建立起“你必须把假说写清楚，我才能判断是否值得资助”的规范。这是一次理性的进步：它把公共资源从空想中回收，投向了可检验的研究。

过程滤网（进展必须可量化）的合理起源在于对“拖延”的纠偏。大型科研项目——从曼哈顿计划到阿波罗计划——教会了管理者一件事：复杂工程需要里程碑，否则可能永远走不到终点。节点审查、阶段性成果、中期汇报，这些工具在工程项目中是高效的管理手段。

出口滤网（成功必须可评价）的合理起源在于问责。当科研经费来自税收或市场竞争，资助方有正当权利问：“我投的钱出了什么？”论文发表、引用率、专利数、市场转化效益，这些指标为此提供了可操作的衡量标准。

这三重滤网分开看，每一条都有正当理由。问题发生在它们的耦合。当入口滤网要求研究者必须将未成形的问题包装成“可陈述的假说”以通过评审，当过程滤网要求研究者必须在预设的时间窗口内交付“可量化的进展”以保住经费，当出口滤网再度以同一套量化指标来判定成功时，这三重滤网就不再是三个独立的合理性工具，而耦合为一个系统性的筛选程序——它的筛选方向是：保留那些能被构成性知识稳定地描述、评价和催生的“创新产出”，筛掉那些需要在不可陈述的混沌中悬置、在无法量化的漂移中试错、从对话性打断中发生断裂、在裁决性肯定中做出无法预先论证的取舍的“可构成者发生条件”。滤网耦合排斥的并不是某一种创新方法，而是让新的认识者得以生成的那些最初始的条件。

4.2 滤网的内在矛盾

然而滤网自身的运作也在积聚瓦解自身的张力。以过程滤网中最典型的量化手段——阶段性成果考核——为例，它在顶尖实验室中催生的系统性“逆向适应”已屡见不鲜：管理者催着要“可展示的成果”，研究者便将手头任何半成品包装成“里程碑交付”。这些交付满足了中期评审，但它们是否为真正的后续突破提供了有效路径，评审本身无法判断——因为评审的标准正是“它看起来像不像一个里程碑”，而不是“这条路最终走不走得通”。长期运作的结果，是过程滤网生产了大量的合规产出，而这些产出对于突破的发生条件——悬置、漂移、打断——的排斥，反倒在于体制内部引发了日益尖锐的抱怨：为什么我们每年产出的“阶段性成果”越来越多，真正改变方向的东西却越来越少？这个抱怨本身，就是三重滤网沉淀下来的内部矛盾，而且是滤网自身无法解决的矛盾，因为要缓解它，就必须让滤网在某个环节松动，而滤网的设计逻辑正是防止松动。

本节不作任何修复方案的宣告，而仅仅记录三件事的同步发生：第一，三重滤网各有合理的发生起源，并非恶意产物；第二，它们的耦合形成了一个结构性排斥方向，排斥的对象恰恰是可构成者发生的必要条件；第三，排斥过程自身也在积聚内部矛盾，这种矛盾是否终将逼出新的滤网形态，不在本文讨论范围之内。但这里至少可以给出一个负向的确认：如果不理解可构成者的发生结构，任何“改进滤网”的方案都将不可避免地用另一套滤网去替代这一套滤网，再次把不可陈述、不可量化、不可预测的环节筛掉——因为管理者不认识那些环节，只能筛掉。

五、对勘两大强对话者：双环学习与类比推理范式

在完成核心案例论证和滤网分析之后，本节正面进入两个必须处理的强韧对话者——它们来自组织学习理论与当代认知科学，分别从不同方向对“可构成者”概念的不可还原性构成威胁。

5.1 阿吉里斯的双环学习：反思何以不触及断裂？

阿吉里斯和舍恩的“双环学习”理论，是组织学习领域中最直指本文核心地带的对话者。当一个组织或个体不仅反思行动策略（单环学习），而且开始反思指导行动策略的那个“主导变量”——那些被默认的假设、价值观和规范——时，双环学习就发生了。这与本文所说的“认识者自身被重新构成”高度重叠。如果双环学习确实是一套可以被培训、被植入组织流程的方法，那么它是否构成了对“可构成

者不可被方法论化”这一主张的直接挑战？

我的回应如下。双环学习所触及的“主导变量”，在性质上不同于可构成者所要切断的那个底层预设。在阿吉里斯的经典案例中，一个管理者通过反思发现，自己行为背后指导行动的主导变量是“我要赢”——这个假设导致他在团队讨论中总是压制不同意见。双环学习让他识别这个假设，并修正它。但这个整个过程都发生在一个更深的、未被触动的框架内部：识别和修正主导变量的方法本身——把假设明确说出来、检验它的有效性、设计新行为——本身依赖于一套构成性知识（如沟通理论、反思技巧、引导流程）。换句话说，双环学习虽然反思了“我为什么要这样做”，但它从未让反思者跌入一个“我连这整套反思方法本身是不是对的那个地基”都不敢确定的深渊。双环学习能够反思的，永远是在给定框架内被“说清楚”的主导变量；而可构成者所发生的地方，恰恰是那个被预设为“可以说清楚”的语言本身在某一点上失去了覆盖力——在那个点上，认识者连“用什么东西去反思”都丧失了。这不是“多反思一层”可以抵达的。双环学习是给定框架内的深度反思；可构成者是框架本身的断裂。两者的门槛是不同的。

5.2 类比推理与概念组合：为什么组合不是发生

当代认知科学对创新思维的主流解释路径之一，是通过类比推理和概念组合来刻画“如何从已知生成新知”。根特纳的结构映射理论论证人们通过将将一个熟悉的源域的关系结构映射到不熟悉的目标域来产生新理解；霍利约克和萨加德的多重约束理论进一步将类比、归纳和演绎整合为一个统一的解题模型。这些研究提出了一个强韧的主张：创新可以被理解是人类认知系统对已有知识表征的重新组合与映射，而并不需要一个神秘的“断裂”概念。

这个范式对于解释大量真实的创新——从仿生设计到商业模式迁移——是极其有力的。但它可以解释爱因斯坦案例中的“可构成者”过程吗？我的回应是：类比推理和概念组合能够解释的是“当新问题已经有了一个可被类比的目标轮廓”时的创造性跳跃。凯库勒梦见蛇咬住自己的尾巴而想到苯环的环形结构，这个跳跃被认知科学视为结构映射的古典范例——蛇咬尾的闭合结构（源域）映射到碳链的可能闭合（目标域）。但请注意：凯库勒在跳跃之前，他的目标域——“苯分子碳链的拓扑”——是被给定知识（化合价理论）已经明确界定了的问题空间。他的跳跃是找到了一个已知源域去填补那个已知问题空间中的一个缺口。

爱因斯坦的情况则不同。1905年之前的物理学共同体并没有一个叫做“时间究竟是不是绝对的”的“问题空间敞口”摆在那里等着人去填空。在洛伦兹和几乎所有同时代人的认知框架中，时间的绝对性不是“一个尚未被解决的开放问题”，而是“不成为问题的不言自明的前提”。可构成者的发生不是“找到了一个更好的类比源域”，而是在“要解决的问题本身是什么”的重新定义中，那个定义过程使得原先不被视为问题的一个预设第一次暴露出它的可质疑性——而这种重新定义，恰恰是在类比推理模型无法预先设定“目标域问题空间”的前提下完成的。换言之，类比推理能够有效运作的前提，是目标域已经被构成性知识界定为“一个问题”；而可构成者要解释的，正是那种“要先把不被视为问题的事情逼成问题”的认知跃迁。这是两个不同层级的过程，后者不能被还原为前者的一个特例。

六、守护可构成者：一个自反性的防守框架

如果前文的论证成立——可构成者的发生依赖于悬置、问题漂移、对话性打断与裁决性肯定这四个不可被方法论化的环节——那么接下来的问题就是：一个组织或共同体，是否有可能辨识和守护这些条件？

这个问题自身是危险的。任何试图把“守护可构成者”转化为一套操作指南的努力，都将立即陷入三重悖论。第一重：辨识悖论。悬置和问题漂移最活跃的阶段，恰恰是它们无法被构成性知识识别的阶段——一旦它们被清晰地识别为“正在发生的可构成者”，它们通常已经被压进了可陈述的模具，不再是原来的那个发生过程。第二重：指标悖论。如果为悬置设立指标（如“允许研究员三年不发表”），管理者将以完成这项指标而行动——这立刻改变了研究员悬置的经验性质：他的悬置从“没有外部压力”变成了“有一个叫悬置的指标在保护他”——两者的差异，相当于“自发沉默”与“被要求沉默”之间的差异。第三重：裁决悖论。裁决性肯定的发生是一条单向路；如果组织试图“保护裁决”，它就必须预先知道哪一个裁决是“真裁决”而非“伪构成”——而做这个判断恰恰是组织自己所做不到的，因为组织的判断工具正是构成性知识。

面对这些悖论，本文不试图给出正面操作规程，而是提出一个自反性的防守姿态：守护可构成者，首要不是去“培养”可构成者，而是去暴露、识别和不断限制那些正在扼杀它的条件。具体而言，可以追问三个问题。

第一，这个组织的构成性知识库存在哪些地方，已经把“什么算一个问题”的边界封死了？封死边界并不总是坏事；在大多数常规活动中，这是效率和可靠性的基础。但当一个组织发现自己在某个领域反复遭遇“预料之外的失败”且这些失败无法被现有框架容纳时，它应该能够识别：那里可能有一个正在积累的边缘失效，而它不应过早地用构成性知识去解释掉那种“说不清的不安”。

第二，这个组织的评价体系是否给“无方向的漂移”留出了空间？注意，这不是说“给漂移设立专项基金”，恰恰相反，任何专项化都重新将它纳入指标。真正有效的守护更接近于不做什么：在某些环节，有意不设置中期审查；在某些节点，不要求用给定术语描述“进展”；在某些档案里，沉默不被解释为失败。

第三，也是最反直觉的：这个组织的对话结构，是否提供了某种形式的“贝索”——不是导师、不是评审者、不是绩效谈话中的上

级，而是一个可以承接混沌、不急于用给定知识去消解它的听众？在许多高密度研究环境中，对话几乎总是功能性的——讨论设计、分析数据、得出结论。而可构成者所需要的对话，往往是松散、漫长、不产出直接结论的。守护这种对话，不意味着要把它制度化；它更需要的是不被挤占。

这三点不是一套完整的操作手册。它们的作用更接近一组提醒：提醒组织其自身的构成性知识库存和评价体系，正在持续地、系统性地排除某些不能被提前识别为“有用”的东西。这些提醒自身也可能蜕化为新教条——比如某天一个管理者把“不设中期审查”变成一个僵硬的规则，对每一个项目都做同等待遇。生成机制的态度不是给出免败的方法，而是持续暴露某种条件自身的不可持续性：今天在组织A里适用的一个防守姿态，明天可能恰恰组织A的整个创新管理文化把这一姿态吸收为绩效展演，它就失效了。守护本身是一个持续被重新追问的问题，不是一个可以被一次性解决的“制度设计”问题。这一节将这种自反性贯穿始终，使防范框架不以“回归本质”的道德呼吁收束。

七、反驳与证伪：使理论自身接受攻击

本节正面回应三个最棘手的反驳，并将可构成者的蜕化条件作为证伪命题。这不是为了显示理论的周全，而是为了在理论最坚硬的内核上开凿可攻击面——可构成者不可被方法论化，不代表它自己不可以被反驳。恰恰相反，如果一个理论找不到自己会失败的明确方式，它就退回了既成事实论——一种自我免疫的教条。

7.1 伪构成的诱惑：表演断裂

第一个反驳来自“伪构成的诱惑”。一个管理者、一个研究生、甚至一个领域的资深学者，在阅读了可构成者的描述后，可能会自觉地、或半自觉地模仿它的叙事——宣称自己经历了某种“断裂”或“裁决”，而实际上这只是用新学到的术语为自己的职业轨迹编织的一套事后叙事。如果伪构成与真可构成者在外部表现上难以区分（两者都可能说“我曾经长期感到不安”，都可能描述“一次关键对话”，都可能声明“我做出了一个无法从现有知识推出的选择”），那么可构成者理论是否无意间为认知自恋提供了一套华丽的说辞？

我的回应分两层。第一，伪构成与真可构成者之间，存在一道生成机制判断标准，虽不能实时鉴定，却可以在时差中暴露。真可构成者的发生过程，会留下不可逆的认知结构痕迹：在断裂发生之后，认识者看问题的底层分类体系被实际改变，不是修辞层面宣布自己被改变了，而是在他此后处理具体问题时，他调用的分类、他所认为的“不言自明的前提”、他绕过陷阱的方式都发生了可追溯的变化。伪构成者则相反：他的“断裂叙事”是自成一体的、自足的故事，但在处理新问题时可能会不经意地退回到构成性知识的标准操作，因为他的底层分类从未真正被动过。这不是不可操作的心理推测；这在长时段中可以由独立于其叙事的第三方——其学生们、合作者、批评者——从其学术行为的实际改变中判断其认知结构的非连续程度。

第二，更重要的一层是：可构成者理论本身并不为这种伪构成提供道德谴责。它不被设计来“检测谁是真正的可构成者”，因此伪构成的存在不是理论的失败。理论的任务仅是揭示：存在一个不能被构成性知识还原的发生过程；至于历史上哪个具体人物是不是真正经历了这一过程，那是传记学问题，不是认识论问题。许多在追忆中讲述自己“断裂”的人可能只是事后叙事，这不妨碍可构成者作为一个分析范畴的有效性——就如许多人声称自己“经历了范式转换”而其实只是在赶时髦，这并不否定库恩的“范式转换”的理论价值。

7.2 AI创新：为什么“重构”不可自指

第二个反驳更根本。随着大语言模型和自修改架构的发展，一个经常被提出的主张是：AI也可以经历某种类似“认知者重构”的过程——它可以通过强化学习的反馈循环调整自己的权重，可以在某种程度上“修改自己的认知框架”。如果这个主张成立，那么将“可构成者”绑定在“认识者有一个自身可以被构成”的前提上，就只是在用生物沙文主义预先排除非人类认知体。

但请认真对待所主张的事情的性质。当一个AI系统进行自修改时，它的修改操作依赖于一个给定的评价函数（哪怕这个函数是从数据中学习的，其学习目标仍然是被人类预设的）。AI的“裁决”从不会被逼到一个没有评价函数的荒漠地带——而可构成者的裁决性肯定恰恰发生在所有既有的评价标准都不再能给出确定答案的那个点上。爱因斯坦在1905年春天的裁决，其裁决性不在于他选择了一个指标更高的选项（在当时的构成性知识框架内，并没有一个现成的函数能把“光速不变”与“绝对时间”放在同一个量纲上比较得分），而在于他选择了改变“什么才算是被比较的对象”的那个分类框架本身。这不仅是“选择了一个更好的选项”，而是“重新定义了计分的东西”——这正是AI在任何已知架构中做不到的事情。AI可以在既有目标函数下调整策略，但它无法在缺乏目标函数的情况下，主动将“目标函数本身应当是什么”作为一个问题从内部逼出来——因为它没有一个在认知压力下会感到“说不通”的自身。它的认知压力不是它自己的，是它的目标函数的设计者的，它仅仅运算那个压力。

7.3 精英主义指控与伦理边界

我必须直面这个最沉重的反驳：可构成者的概念是否在无意间为一种关于创新的社会达尔文主义背书？它描述的过程——悬置十年、

无关KPI、来自边缘的认知断裂——听上去只适用于爱因斯坦这种天才。对于在实验室里、在设计桌前、在讲台上每天面对认知挣扎的普通人，这个概念是否过于冷酷了？它是否暗示着，如果你的挣扎没有产出断裂性的重构，那只是因为你的“裁决性肯定”尚且不够、你的“压力地形”尚且不是天选——从而将结构性困境转译为个体的认知资质缺陷？

这个反驳切中了概念的伦理要害。让我尽可能诚实地回应。

首先，可构成者被定义为事件，而非资质。它描述的是一个认知过程在特定条件组合下实际发生了的事情，而不是一种“拥有”或“不拥有”的内在禀赋。天才是事后追认的神话——我们把那些发生了可构成者并被历史记住的人，命名为“天才”；但这并不意味着在发生之前，他们就已经是天才。十六岁的爱因斯坦没有任何理由被当时的任何人预言为“将成为可构成者”。把可构成者当作天才的勋章，是搞反了发生顺序：不是天才导致了可构成者，是可构成者在事后被误解为天才的证据。

其次，可构成者的概念因此对“普通认识者”有一个本应该说但经常没有说出来的话：它描述的不是认知等级，而是认知条件的极值。每一个在边缘处挣扎的认识者，都可能在某个没有被现有知识覆盖的缝隙里经历边缘失效——区别不在于“你是不是可构成者的料”，而在于你所处的结构是否允许那个缝隙不被过早地堵上。爱因斯坦所享用的“十年悬置”并不是一种心理能力，而是一种社会偶然——专利局的边缘位置、没有发论文的KPI、没有必须取悦的基金评审人。从概念上说，可构成者理论恰恰要暴露：大多数普通认识者的挣扎之所以没有演变为断裂，不是因为他们缺少某种天赐的资质，而是因为他们所处的评价结构从来没有给他们的“说不清的不安”留出足够长的不被消解的时间。在一个要求每位博士生在第三年必须产出可发表的阶段性成果的环境里，即使有一个十六岁的爱因斯坦坐在那里，他的追光疑惑也将在第一学期被导师友善地解释掉：“孩子，没人跑得到光速，别瞎想了，去做点正经课题。”

这引出了伦理层面的一条底线。可构成者的概念不能被用来做筛选——“我辨识出你有可构成者的体质，所以给你特殊投资”。所有试图从概念中提取识别指标的做法，都恰好背叛了这个概念的出发点：悬置之所以成为悬置，恰恰因为它没有被当成悬置来特别对待。但概念可以被用于做另一种判断：当我们在管理创新、设计评价体系、指导年轻研究者时，至少不应该让我们自己的管理工具，去主动消灭那些我们根本无法识别的发生条件。这个概念给出的不是“谁能创新”的答案，而是一个更谦卑的警示：我们现有的评价和管理体系，正在系统性地清除它所不能识别的东西——而我们至今不知道我们因此已经筛掉了多少可能的可构成者。

7.4 证伪条件：让理论自身可以失败

一个生成机制的理论必须敢于给出自己的失败条件。以下命题构成可构成者理论的可证伪核心。

第一，可构成者的非对称发生命题：如果在历史上能找到一例情况，其中认知者A与认知者B面临同样一组反常数据、共享同样一套构成性知识、拥有同样长期的悬置条件、同样经历了与不共享底层预设的他者的对话性打断——且这两人的认知构成史并无决定性的不同，但A发生了断裂而B没有发生，并且这个差异不能通过“裁决性肯定的压力地形差异”来解释，那么本文关于可构成者发生机制的普遍性主张即被削弱。

第二，对话性打断的结构必要性命题：如果存在一个确凿的历史或实验案例，证明一位认知者在完全没有任何他者对话性打断的情况下——即其整个断裂过程发生在其内部独白中、从未被“对另一个特定主体说清自己的混沌”的那个压力所逼——仍然完成了同样性质的底层认知框架的不可逆断裂，那么本文关于“对话性打断是发生结构中的必要环节”的主张就需要被修正。

第三，裁决性肯定的非任意性命题：如果存在一个案例，其中的裁决确实仅凭任意的个人偏好就完成了突破，而无法被追溯到一个由长期认知压力所形成的特定压力地形结构，那么本文关于裁决性肯定的认识论刻画就退化为单纯的心理学叙事。

这些证伪条件本身不应被理解为本理论的“安全阀”——即“如果出现反例，我会用这一条去解释掉它”。它们被摆在这里，是为任何愿意认真检验这一理论的批评者提供攻击面。这可能是本文能为自己做的最诚实的一件事。

结语：知道者自身的发生与不安全

本文始于一个追问：当我们把创新当作一种“特殊的知识产出”来管理时，我们是否忘了追问那个产出者在被摆到“创新者”这个位置上之前，他自身是如何从知识中被构成的？为了回答这个问题，本文提出了“构成性知识”与“可构成者”的不可还原区分，并通过爱因斯坦与洛伦兹的对称解剖，展示了可构成者发生的非必然结构。

但诚实地说，整个研究所做的，只是把这个追问推到了它目前能被推到的位置。它在三个方向上仍然是不安全的。第一，可构成者的发生结构中的四个环节——悬置与边缘失效、问题漂移、对话性打断、裁决性肯定——目前只在一个核心案例和一个对称失败案例中得到了充分展示。它尚未被扩展到其他学科和领域的突破性创新中接受检验，这是历史性的不饱和。第二，裁决性肯定中的“压力地形”概念虽然在本轮修订中被推向了认识论层面，但它仍然不是一个可被独立测量或建模的变量——它目前的解释力是回溯性的，是“在知道结果之后再去追溯认知者的认知构成史如何塑造了他的裁决”。没有向前预测的能力，就没有最终的认识论完满。第三，也是最根本的：本文所防范的那种“守护可构成者退化为新教条”的风险，不止是本文所讨论的对象的一个外部威胁，而是本文自身也可能制造的风险——也就是说，如果一个读者读完本文后，觉得自己“懂了可构成者”，从而可以用这个概念去“解释”一个又一个的突破案例，去在讨论班上

向别人讲述“你看，爱因斯坦就是经历了裁决性肯定”，这个概念就已经退化为一种新的构成性知识标本。它重新回到了它本来要切开的
那种“认识者不变、知识在手中流转”的图景。在那幅图景里，说话的是一个稳坐于可构成者概念之上的解释者——而这个解释者，恰恰
没有被重构过。

因此，本文没有一句话可以保证自己不会蜕化。它最后的诚实，是把蜕化的若干路径写进第七节的证伪条件，让这些条件也成为可以
被检验和攻击的对象。可构成者这个名词是否能逼读者去追问自己身上的构成史，而不是仅仅被当作一个分析别人家突破的工具——这是
本文留给读者自己去判断的事。这个判断，无关方法论，无法被纳入框架的收束。它是“谁”的问题。

深化增补：被放跑的最近邻、操作化与证伪

以下四节为编辑增补，与作者正文分开标注。

增补一：与福柯“自我技术”、达斯顿-加里森“科学自我”的硬边界

正文与库恩的格式塔转换做了硬论证，又对勘了双环学习与类比推理范式，却放过了专门研究“认识者如何被塑造”的那一支
传统。福柯晚期的自我技术，论证主体通过一整套实践把自己构成为特定类型的主体；达斯顿与加里森的科学自我，则具体考察了
不同时代的科学德性如何要求并锻造出不同类型的观察者。两者都在做本文所称的“追问产出者如何被构成”，不划界则可构成者
会被读成这一传统的个案应用。分野在构成的施动方向。自我技术与科学自我描述的是一套已经存在的规范如何把认识者塑造成合
乎该规范的样子，构成的模板先于被构成者；本文的可构成者恰恰发生在没有模板的位置上——认识者被逼到既有构成性知识无法覆
盖的边缘，重组之后成为的那个人，不合乎任何在先在的规范类型。可判别面由此清晰：若一切重大的认知重组都能被追溯到某种在
先在的德性规训或自我实践，则本文没有独立地位；洛伦兹与爱因斯坦的对称解剖正是为此设计的——两人共享同一套科学德性与训练
传统，结果却分岔。

增补二：四环节的可判别形式

悬置与边缘失效、问题漂移、对话性打断、裁决性肯定四环节，需要可判别形式才能支持文末的证伪命题。建议：边缘失效，
指认识者明确指认既有工具在某一具体问题上不是用得不够好，而是不适用；问题漂移，指其追问对象从原问题移到该原问题所依
赖的底层预设上，可由前后文本中问题句的宾语变化辨识；对话性打断，指存在一个具体他者，其提问使认识者无法沿原路继续，
且该打断不可由认识者自问自答替代——这一条最关键，也最容易被事后美化，因此须以同时代通信、笔记等旁证为准，不采信回忆
录；裁决性肯定，指认识者在证据尚不充分时对新框架做出承担性肯定，可由其此后放弃的备选方案数量近似。四项须全部出现方
计为一次完整事件，缺一即计为未完成——洛伦兹一侧正应据此被记录为在第二与第三环节处停住。

增补三：三重滤网的操作化指标

正文追溯了当代创新管理中三重滤网的形成史，并指出滤网自身的矛盾正在积聚。若不给指标，这一诊断无法与常见的评价体
制批评相区分。建议三项：悬置容忍度，即一个单位允许一名研究者在无阶段性产出的情况下持续工作的最长时长；漂移容忍度，
即课题目标在执行期内被实质改写而仍可继续获得资源的比例；肯定成本，即一个尚无充分证据的方向从提出到获得首笔资源所需
经过的评审环节数。三项都不测量产出，也不涉及任何创新水平的判断，因此可用于横向比较不同单位而不引入绩效话语——这一点
本身是本文的实质主张：凡以产出为测量对象的指标，都只能看见构成性知识那一侧。

增补四：与作者已发表《探究的会计化退化》的分野

作者已发表的《探究的会计化退化：当代知识生产评价机制与认知主体性的腐蚀》与本篇共享同一诊断对象——评价机制对认知
主体的损害——必须切开。前篇的分析单位是评价机制本身：它追踪会计化如何一步步接管学术判断，受损者是整个知识生产场域，
其证据在制度与指标的演变上。本篇的分析单位是单个认识者的重构事件：它追问在压力之下那个能看见新问题的人如何被构成，
以及为什么滤网会系统性地筛掉这一过程，其证据在具体个案的四环节是否完整上。两篇的关系是诊断与机制：前篇说明场域被什
么侵蚀，本篇说明被侵蚀掉的究竟是哪一个环节。作者宜在本篇明确引用前篇，并指出本篇的三项滤网指标可作为前篇诊断的可操
作接口。

编辑增补所依据的核验文献

Michel Foucault, *Technologies of the Self*, University of Massachusetts Press, 1988.

Lorraine Daston & Peter Galison, *Objectivity*, Zone Books, 2007.

Gerald Holton, *Thematic Origins of Scientific Thought: Kepler to Einstein*, Harvard University Press, 1973.

Olivier Darrigol, *Electrodynamics from Ampère to Einstein*, Oxford University Press, 2000.

Chris Argyris & Donald A. Schön, *Organizational Learning: A Theory of Action Perspective*, Addison-Wesley, 1978.

Thomas S. Kuhn, *The Structure of Scientific Revolutions*, University of Chicago Press, 1962.

Gilbert Ryle, *The Concept of Mind*, Hutchinson, 1949.