

在平滑的刺痛中：AI如何静默替换创作者的内部校准

——不是能力在衰减，是另一条更快的回路在生成

王德生 · 德麦国际 · 2026年7月17日

摘要

本文论证，在生成式AI深度嵌入创作过程的当下，一个此前未被精确把握的转变正在发生。设想一个写作者，冬夜独行时被某扇窗户里的灯光轻刺了一下——不是感动，不是怀旧，只是一种说不清楚的、沉在身体层面的注意力偏斜。此后数日，她没有试图去“抓住”它，而是让这个感受松散地吸附一些偶然读到的材料：玻璃工艺史中关于手工玻璃不规则气泡的记载，一本讲城市夜行动物的科普。她开始写一些试探性的片段，并依据重读时身体的微小反应——这个句子“身上不对劲”，那个句子“好像离什么近了一点”——来决定保留或删除。在反复划掉与重写之后，某个下午她突然意识到，她一直在写的不是灯光，而是距离：一种“隔著玻璃看见却不可能触及”的距离。这个语义内核不是被“想到”的，而是在反复的身体校准中被**凝定出来的**。本文将这条从身体标记浮现、经由松散吸附、跨域试探与身体校准的反复练习、最终推进为清晰语义结构的过程，命名为**内部校准链条**。在此基础上，本文的核心工作是：追踪并解剖当AI的即时平滑化介入链条最关键一环——身体校准的反复试探——时，创作者的身体感受发生了何种质地变化，由此生成了何种新的感知模式与创作姿态。本文论证，被“润色”截断的那一瞬间，不是一次简单的“省去麻烦”，而是一种身体感受的质地转变：从向外试探的刺痛，被平滑为一种向内塌缩的轻微挫败。创作者由此被训练出一套新的内部动作——不是训练如何推进自己的模糊感受，而是训练如何在看到平滑版本后，迅速调整自己的感受以与之匹配。这不是“校准能力的衰减”，而是一种**替代性校准方式的发生**。本文将这一新方式的生成机制定位为“执行者的位移”——传统竞争催促创作者自己跑，AI在创作者尚未决定要不要跑的时候就开始替他动。这一转变的后果之一，是创作者在作品重读时那种从身体内部升起的、知道“这一步踏对了”的刺痛感正在被一种更温和的、“这看起来不错”的合理感所替代。本文回应了“流畅工具无罪论”与“高阶专家直觉”两个反驳，并在承认技术已不可逆的前提下，指出未来的实践不在于“保护”旧的校准方式，而在于能否设计出新的技术与制度，来重新激活创作过程中那种不可被外部标准化替代的、在身体内部反复凝定的张力。

关键词：内部校准；生成式AI；身体感受；替代性校准；具身认知

一、引言：一个正在发生但尚未被命名的转变

关于AI与创造的讨论，在过去几年间几乎演变成一场全民参与的世纪论辩。视觉艺术家起诉AI公司未经许可抓取作品用于训练，作家们联署公开信抗议大语言模型对文字行业的侵蚀，音乐学院开始修改毕业论文的原创性审查标准，立法者则在“AI生成内容能否享有版权”这一问题的技术迷雾中进退失据。这些喧嚣背后的核心焦虑，可以用一个朴素的问题来概括：当机器也能写出像样的文章、画出精美的图像、谱写悦耳的旋律时，那个叫“原创”的东西，还剩下什么？

现有的回应大致分布在两条路径上。第一条追问的是归属与权益：AI产出的内容算不算“作品”？如果算，著作权归谁？训练数据的使用是否构成侵权？它关涉的是创作者能否以及如何AI时代继续以创作为生，但始终在回答一个更表层的问题——蛋糕怎么分——而没有触及那个更深层的问题：蛋糕本身是否在发生某种质地上的变化。第二条路径追问的是能力与本质：AI有没有“真正的创造力”？人类创造力的独特之处究竟何在？这些讨论在哲学上富有启发，但在实践中却常常悬在半空——它们解释了人类为何独特，却没有说明当AI深度嵌入创作流程后，那种独特是否正在被系统性地重新配置。

本文的工作始于对这两条路径的双重不满足。在“分蛋糕”与“论本质”之间，存在著一个更隐蔽的地带。这个地带不关心AI是否“真正”具备创造力，也不关心版权归属的规则如何拟定。它关心一个更微观的问题：当创作者开始使用以“最小化预测误差”为训练目标的AI工具时，在创作者的身体与机器之间每一次微观互动的那一秒钟内，**实际发生著什么**？这个地带尚未被充分照亮，因为它要求我们悬置对“创造力”、“想像力”、“灵感”这些大词的巨观讨论，转而进入一个极其微小的、几乎不可见的层面——那个在创作者身体内部反复校准、试探、推进的过程。

本文的核心工作是：追踪并解剖这一微观过程，并在此基础上论证，当前正在发生的转变，不是“创作者变懒了”或“原创性衰减了”——这类诊断本身就是需要被悬置的——而是一种**替代性校准方式的发生**。AI工具的即时平滑化，在创作者尚未完成身体校准之前，就给出了一个平滑的替代方案。这不是对一条既有链条的“截断”，而是对一种新的感知-行动回路的反复强化。在这个反复的过程中，创作者的身体标记仍然会浮现，但那条将标记推进为清晰语义结构的旧有链条，正被一条更短、更平滑的新链条所覆盖。由此发生的，不是一个“更少”的创作主体，而是一个**不同的**创作主体——其内部校准所依赖的身体敏感性，正在从“向内试探的刺痛”被重塑为“对外部平滑选项的快速识别与接纳”。

这一定位，将本文与两种常见的立场区分开来。第一种立场是技术怀旧主义——它哀悼一个正在逝去的、前AI的创作黄金时代。本文不共享这种乡愁，因为这项追问的起点恰恰是：从来不曾存在过一个“未经任何技术中介”的纯粹创作状态。从口头传统到文字书写，从纸笔到打字机，每一次媒介技术变革，都在重新配置创作者与自己身体感受之间的关系。问题从来不是“配置”本身，而是当前这一轮配置的**特定方向与后果**。第二种立场是技术乐观主义——它声称AI只是更好的工具，有自律的创作者完全可以选择性使用，不必担心任何能力的削弱。本文将在第六节正面回应这一立场，但在此先指出：这一立场预设了一个完全自主的、能够清晰区分“使用工具”与“保持能力”的理性主体，而创作过程中最关键的校准恰恰发生在这一主体能够清醒决策之前的前反思层面。

在正面展开论证之前，需要先回答一个前置的追问：凭什么说这个问题“尚未被命名”？认知科学中有一个历史悠久的研究传统——从Ericsson等人（1993）的刻意练习理论，到Anderson（2010）关于认知技术如何重组人类技能的神经可塑性研究——已经反复论证过一个命题：当一个外部系统接管了某个认知环节后，相关的神经通路会发生使用依赖性的重组，不用的能力会衰退。这个“用进废退”的道理并不新鲜。如果本文的工作只是把这个道理搬到创作领域再讲一遍，那它就没有提出任何新东西。

本文的论证恰恰始于对这一传统的追问。它说对了一件事——“用进废退”，但它没有回答一个更关键的前置问题：在AI与创作的具体互动中，**那条被废用的链条的具体环节是什么**？当我们说“用进废退”时，我们究竟在说什么“用”和什么“废”？是“写作速度”吗？显然不是，AI恰恰让写作更快了。是“语法能力”吗？也不是，AI提供的正是语法上更通顺的版本。是“想像力”吗？AI甚至可以生成令人惊艳的新奇组合。那么，那个正在被废用的东西，究竟是什么？它发生在创作过程的哪一个具体环节？被废用之后，创作者身体里生成了什么新的东西来填补这个空缺？这些问题——被废用的**那一条链条的具体环节**，以及由此发生的**新的替代性校准方式**——才是本文所要填充的空白。它不是对“用进废退”的重复，而是对这一命题在创作领域的机制级精确化：定位那条链条，解剖它的环节，并追踪一条新链条如何在旧链条被绕过的过程中，被反复强化为一种新的创作姿态。

本文将沿著以下路径展开论证。第二节诊断AI训练逻辑中“最小冗余”原则的形成，以及它如何从一项技术指标滑向一种价值判断，重点揭示这一滑移如何作用于创作者的身体感受与创作节奏。第三节通过一个创作案例的逐环解剖，追溯一条从模糊身体感受到清晰语义结构的完整推进链；该节同时包含了内部校准与深感／聚焦传统、刻意练习／专家直觉等既有理论的严肃对话，以确立这一命名的不可还原性。第四节则对整条链条中最关键的一环——身体校准的反复试探被AI即时平滑化介入时——进行显微镜级别的机制描述，重点解剖创作者身体感受在那一刻的质地变化，并在此基础上命名一种“替代性校准”方式的发生。第五节简要分析这一转变在创作者个体、作品品质与创作生态三个层面的后果。第六节回应两个最严肃的反驳：“流畅工具无罪论”与“高阶专家直觉可覆盖论”。第七节将论证推衍到教育与学术研究领域，论证内部校准的替代性发生并非AI时代所独有，而是任何“外部平滑机制介入微观创作过程”的场域都可能出现的现象。第八节总结全文，在承认技术已不可逆的前提下，指出未来的实践不在于“保护”旧的校准方式，而在于能否在承认当前这一替代性发生的基础上，重新设计技术与制度，以激活创作过程中那种不可被外部标准化替代的、在身体内部反复凝定的张力。

二、从平滑到替代：AI如何介入创作过程的内部校准

要理解AI如何重新配置创作者的内部校准，不能只停留在描述其训练原理或批判其商业模式。必须进入一个更微观的层面：创作者在使用AI工具时，身体与机器之间实际发生了什么。

（一）AI的“平滑本能”：从损失函数到创作界面

当前主流的生成式AI模型——无论是大语言模型还是文本到图像的扩散模型——尽管架构各异，却共享一个根本性的训练目标：在给定输入条件的情况下，使模型的输出与“正确”输出之间的差异尽可能小。在技术术语中，这个差异被称作“损失”。训练过程，就是持续调整模型内部的数千亿个参数，让这个损失函数的值不断下降。从信息论的视角看，损失最小化的过程，等价于压缩数据中的冗余——消除那些“不必要的”、“额外的”、“无法帮助做出更准确预测”的信息。

这个逻辑在它自身的工程语境中是极其合理的。一个预测下一个词的模型，如果总是能精准命中目标词，它就是一个完美的模型；如果它经常“说错话”，工程师就会认为它的损失还不够小，需要继续优化。问题在于，“精准”和“正确”在这里的评判标准，并非基于真理或创意，而是基于与训练数据中已有模式的吻合度。统计上最平均、最不意外的输出，意味著最低的损失。统计上的离群值——那些罕见、怪异、不合常规的表达——则意味著高损失，是模型要努力消灭的对象。模型被训练得越好，它的输出就越平滑、越可预测、越趋近于一种跨语境的“平均美感”和“平均合理性”。

如果这项技术只是待在实验室里，这个逻辑不会构成任何文化影响。但一旦它被封装进消费级产品并投入广泛使用，其影响便沿著三条具体的通道逐步渗透。

第一条通道是**反馈循环**。用户最初使用AI辅助创作时，往往带著试探的心态。他们输入一个粗糙的开头或一个模糊的创意，模型迅速返回几个看起来“还不错”的版本。用户在多个版本中选择、修改、拼接，每一次选择都在强化某种标准：那些语法更通顺、逻辑更连贯、表达更“流畅”的版本更容易被采纳；而那些含有意外联想、不对称比喻或微妙反逻辑的生成物，则被划走、删除或标注为“不理想”。这看似只是一个效率工具的使用过程，但实际上，每一次这样的微观决策都是一次价值判断的演练。“好”的体验逐渐与“平滑”、“合理”、“可预期”绑定，“差”的体验则与“奇怪”、“偏离”、“不可理解”挂钩。创作者未必意识到，自己在一次又一次的反馈中，正在被训练成一台更擅长识别和淘汰“不平滑”输出的“人工筛选器”。

第二条通道是**界面设计**。AI创作工具的界面并非价值中立的。工具栏上“优化”、“润色”、“完善”、“提升”等按钮的存在，本身就是在向用户暗示一个前提：你的初稿是需要被“优化”的，而“优化”的方向就是向更标准、更通顺、更符合规范的方向靠拢。一个写作者打开AI辅助写作软件，看到自己的句子被划上了蓝色的建议线，点击一下，句子立刻变得“更好”了——更流畅、更规范、更像“应该有的样子”。这种体验是即时的、可感知的、带有轻微愉悦的。它不像一个老师站在身后批评你，而更像一个助手温柔地为你省去了麻烦。正是这种“温柔”，使得价值原则的内化比任何强制都更深入——创作者不是被迫放弃自己的偏离，而是被教会了“偏离本来就不值得留恋”。

第三条通道是**使用频率与习惯养成**。当AI工具从一个偶尔使用的“外挂”变成每日写作的默认界面时，它所携带的逻辑就从“辅助”升级为“环境”。一个每天用AI辅助写作的创作者，他的认知流程会发生结构性变化：以前，从模糊念头到成型文本，中间有漫长的酝酿、试错、反复推翻的过程；现在，模糊念头第一步通向的是“让AI先出几版看看”。这不是简单的流程缩短。酝酿期的功能——让远距离的神经回路发生随机联结，让无意识的材料慢慢浮上意识——被跳过了。不是因为AI禁止了酝酿，而是因为酝酿需要“暂时不得到答案”的耐心，而AI的即时反馈让“不得到答案”变得难以忍受。

（二）问题的深化：从“冗余被消灭”到“替代性校准的发生”

现有的批评——包括一些非常有洞察力的讨论——往往将上述过程的后果定位为“创造所需的冗余空间被压缩了”。这个判断是正确的，但还不够精确。它把“冗余”描绘成一种需要被保护的空间或资源，仿佛问题只是“创作者没有足够的时间去徘徊和试错”。这个框架固然成立，但它没有指出一个更根本的微观机制：那些在徘徊和试错中实际发生了什么？当“冗余空间”被AI的即时反馈填满之后，创作者的身体里生出了什么新的东西来填补那个原本由“不确定性”所占据的位置？

本文的论证将在此处推进关键一步。本文主张，创作过程中的徘徊、试错和无目的探索，并非只是“为最终成果做的准备工作”——它们是一个创作者**维持和训练自己内部校准能力的必要练习**。所谓内部校准，指的是一种难以被形式化、却在创作中起著根本性作用的能力：创作者用自己的身体感受——一种轻微的刺痛、一种说不清的抵触、一种隐约的“对了”或“不对”——来判断创作过程中每一个微观步骤的取舍。这种能力不是一种天赋，而是一种需要被反复练习才能维持的技艺，就像一个钢琴调音师必须持续地用自己的耳朵去听那些微小的音差，才能保持听觉对音准的敏感度。

当AI工具提供即时的平滑替代方案时，它所做的，不仅仅是“省去了试错的麻烦”。它在每一次互动中，都提供了一条绕过内部校准的捷径。创作者不需要再去反复感受那个粗糙的句子是否“身上有刺”——AI已经告诉他一个没有刺的版本是什么样的一次、两次、十次、一百次——每一次的绕过，不只是对内部校准能力的一次废用，更关键的是，它在反复地**强化另一条路径**：身体标记浮现→打开AI→看到平滑版本→接受。这不是能力的“空白”，而是能力的**替代**——一种新的内部动作正在被训练出来：不是训练如何推进自己的模糊感受，而是训练如何在看到平滑版本后，迅速调整自己的感受以与之匹配。创作者正在从一个“向内试探的人”，被重塑为一个“对外部平滑选项进行快速识别与接纳的人”。本文将在第四节，通过对一个微观瞬间的逐环解剖，完整地展示这一替代性校准是如何在一次点击中发生、并在反复点击中被强化的。

三、内部校准：从身体感受到语义结构的推进链

“内部校准”是本文的核心命名。但它的合法性，不取决于定义的精确，而取决于它能否从一个真实的创作过程中**被生成出来**。本节选取一个微观的写作案例，通过逐步解剖一个完整的创作过程，展示一条从模糊的身体感受、经过反复试探与自我校准、最终推进为清晰语义结构的链条。在此基础上，本节将与认知科学中两个最相关的理论传统——刻意练习／专家直觉研究、以及创造力的生成-探索模型——进行严肃对话，以确立“内部校准”的不可还原性。[2]

在开始解剖之前，需要先做一个术语上的简要约定。本文交替使用“身体感受”、“身体标记”、“刺痛”三个词。它们的

关系是这样的：身体感受是内部校准运作时创作者所征询的整个感觉场——包括但不限于那些被明确识别出来的前语言信号；身体标记是那一类前语言信号的总称——那些尚未被意识清晰捕捉、但已经在身体层面引起微小的注意力偏斜的冲动或张力；刺痛是本文案例中身体标记的具体质地（其他案例中可能是其他的质地——闷、痒、轻微的恶心、某种说不清的“被拉扯感”）。[6] 本文使用“身体标记”而非Damasio的“躯体标志”，是因为后者专指决策过程中将情绪信号与预期结果联结起来的身体状态，而本文的身体标记则指向更为原初的、在意识边缘涌现的前语言冲动。两者在本体论层次上有别，本文不进入Damasio理论的详细框架。

（一）案例解剖：一个短篇小说的生成过程

写作者C在一个冬天的傍晚，独自走在回家的路上。天已经全黑，路旁住宅楼的窗户里透出暖黄色的灯光。在某个无法确认具体原因的瞬间，C感受到一种说不清楚的触动——似乎关于“窗户里的灯光”这件事有什么东西在隐隐地抓著她，但她无法说出那是什么。不是感动，不是怀旧，不是孤独。如果非要描述，那只是一种细微的、带有轻微刺痛的注意力偏移：好像某个念头在意识的边缘轻轻碰了她一下，随即消失。[1]

这是链条的第一环——**身体标记的浮现**。C的这个感受不是一个可以被清晰陈述的“想法”，而是一种前语言的、沉在身体层面的信号。它没有名字，没有逻辑，没有可以被有效传达给他人的内容。但它不是纯粹的噪音。它有一个极其微弱的指向性——它把C的注意力稍微拉向“灯光”这个方向，虽然还不知道这个方向通往哪里。

在接下来的几天里，C没有试图去“抓住”这个感受。她没有打开文档，没有构思任何情节，甚至没有告诉自己“我在为写作做准备”。她只是在通勤的路上偶尔想起那个瞬间，让它在心里待一会儿，然后放掉。她也翻了几本与此无关的书——一本关于城市夜行性动物的科普，一本讲玻璃工艺史的小册子，还有一篇关于城市下水道系统的长篇报导。这些阅读没有明确目的，只是恰好手边有，恰好有兴趣。那篇下水道的报导，她读的时候觉得“有意思”，但说不清与那个冬夜的刺痛有任何关联——她只是放任自己读完了，然后就放下了。

这是链条的第二环——**松散吸附**。那个模糊的刺痛像一个微弱的引力源，在C的日常活动中吸附一些与之有共振可能的材料。玻璃工艺史中关于早期手工玻璃含有微小气泡、在光线下呈现不规则折射的细节，因为与“灯光”的微弱关联，被吸附过来。夜行性动物对光的特殊敏感，也被吸附过来。那篇下水道的报导——它讲述的是工人们如何在城市地下深处、在人们看不见的地方维持整个城市的正常运转——虽然没有被用在最终的作品中，却在吸附的时刻为C提供了一种“被隔绝的、不可见的劳动”的身体质感，这个质感后来渗进了小说中巡夜人这个角色的身体经验里。这些材料不是被“选中”的，而是被“共振吸引”的——C没有在搜寻它们，它们自己浮了上来。关键在于，有些吸附材料最终会被推进为语义结晶的一部分，有些则只是在某一个阶段提供了必要的共振、然后被放下——但它们在吸附阶段的存在本身，维持了那个最初刺痛作为引力源的活跃性。

有一天晚上，C在洗澡的时候，突然想起玻璃工艺史里读到的那个细节——不规则的折射。不知道为什么，这个细节与那个“窗户里的灯光”的瞬间在脑子里碰到了一起。C仍然不知道这意味着什么，但她隐约开始觉得，“不规则的折射”和“被看见但无法被接近的暖光”之间存在某种她还没有理解的联系。

这是链条的第三环——**跨域的试探性联结**。两个来自不同领域、没有逻辑关联的材料，因为在身体感受层面的某种共振，在C的脑子里发生了碰撞。这个碰撞不是推理的结果，也不能在事前被判断为“有用”。它只是发生了，因为C一直保持著那个感受的活跃，让它在在一个低强度的状态下持续存在。

此后，C开始做一些试探性的写作。她写下几个片段——有的是第一人称的内心独白，描述一个人每晚步行穿过城市，观察那些亮著灯的窗户；有的是第三人称的场景描写，聚焦于一扇特定窗户的内部；有的是完全实验性的文字，混合了玻璃工艺的术语和情绪描述的碎片。这些片段互相之间没有明确的逻辑联系，有的被她保存下来，有的第二天早上看觉得毫无意义就删掉了。这个阶段持续了一周多。关键的是，C在判断一个片段“留还是不留”的时候，依据的不是任何外在的写作规范——不是语法、不是结构、不是可读性——而是她重读时身体的微小反应。有些片段读起来“身上不对劲”，有些片段读起来“好像离那个东西近了一点”。这个判断无法被形式化，但它不是随机的。它是一个创作者长期训练出来的内部校准能力在运作：身体知道什么时候“碰到了”，即使意识还说不清楚“碰到了什么”。

这是链条的第四环——**身体校准的反复试探**。C写下片段，身体感受给出反馈，她根据这个反馈决定保留或删除，然后再写新的片段。每一轮试探都在微调那个最初的引力源本身——感受在被表达的尝试中，被推进得更清晰了一点。这不是“找到了一个事先存在的东西”，而是“感受在试探中被自己重塑为一个更清晰的结构”。内部校准在这里不是一个判断工具，而是一个生产装置——它不仅判断对错，它还在每一次判断中，把那个被判断的对象推进到下一个状态。

然后，某个下午，在反复重读那些片段并划掉大量内容之后，C突然看清楚了一件事。她发现自己一直在写的那些“隔著玻璃看见灯光”、“不规则的折射”、“观察但不进入”的意象，共享一个潜在的结构：一个主体与一个温暖的对象之间，隔著一层无法穿越的透明屏障。这个结构不是关于灯光，不是关于玻璃，而是关于**距离**——一种“被暖光包围却无法进入”的距离，一种“隔著玻璃看见而不可能触及”的距离。这个“距离感”是她一开始完全没有意识到的，却一直在暗中组织著她所有的试探。它不是在试探过程中被“添加”进去的，而是在试探过程中被**凝定出来的**——那些被划掉的片段不是被“淘汰”的废料，而是这个凝定过程中必要的代价。每一次身体校准说“不对”，都在把那个模糊的引力源往“对”的方向推进了一点。当推进累积到一定程度时，引力源露出了它的结构——它一直是一个关于距离的矛盾体：“渴望接近”与“不可能接近”之间的张力。而“距离”这个词，就是这个矛盾的暂时凝定。

这是链条的第五环——**语义结晶**。模糊的身体感受，经过松散吸附、跨域试探、反复校准，最终被推进为一个可以被明确说出、并成为作品组织中心的语义结构。C随后完成的短篇小说，讲述的是一个城市巡夜人每晚记录那些亮著灯的窗户里传出的声音，他从未进入其中任何一间，却在经年累月的记录中发展出了一套独特的、只属于他自己的“遥距亲密”的语言。小说的质地——那种冷静中带著轻微灼痛的叙事语气，那种反复出现的玻璃和光的意象，那种对“不可进入性”的近乎执念的书写——都不是C在出发时能够设计出来的。它们是从那条被走完的推进链中，被生成出来的。

（二）内部校准与专家直觉：不可还原性的核心论证

到此为止，本文已经展示了一条完整的内部校准链条的运作。但在它被确立为一个新的理论命名之前，必须正面回应一个最致命的挑战：这个概念与认知科学中“刻意练习”与“专家直觉”的文献之间，究竟是什么关系？

根据Ericsson等人（1993）关于专家技能获得的研究，在各个领域——从国际象棋到小提琴演奏——顶尖专家的一个标志性特征，是他们在长期刻意练习之后发展出了一种对微观品质的快速、自动化的前意识判断力：棋手一眼看出棋局的关键，小提琴家一听就知道音准的细微偏差，无需经过缓慢的推理。这种“专家直觉”，听起来与本文所描述的“内部校准”——创作者依靠身体的微小反应来判别创作取舍——高度重叠。不仅如此，一个更尖锐的版本是：一个顶尖的理论物理学家在审视一个全新的异常数据时，那种“这里有东西”的强烈直觉，难道不是在标准空缺的情况下，向自己身体化的“理论品味”征询方向吗？如果这就是“内部校准”，那它不过是“高阶专家直觉”在创作领域的应用，为什么需要一个新名字？

本文的回应是：专家直觉与内部校准的根本差异，不在于它们运作的速度或意识参与的程度，而在于**被校准对象的本体论地位**，以及由此决定的**校准动作的性质**。这个差异可以在三个层面上展开，并在最关键的第三层上，论证内部校准的不可还原性。

第一，**标准的性质**。专家直觉的所有经典研究——无论是棋手的模式识别、小提琴家的音准判断、还是物理学家对异常数据的敏感——都共享一个前提：存在一个可以被经验世界所确认的、具有公共性的标准。棋局有输赢（即使最佳走法有争议，胜负是最终的公共判定），音准有正确与否（可以被仪器测量），物理学家的直觉最终需要被实验数据所验证。正是这个公共标准的存在，使得长期刻意练习成为可能——专家之所以能发展出直觉，是因为每一次练习最终都有明确的、来自外部世界的反馈信号告诉他们“对了”还是“错了”。内部校准所面对的，恰恰是这样一个公共标准**尚不存在的**环节。当C感受到那个关于“灯光”的刺痛时，她无法以任何公共标准来判断这个感受是“对的”还是“错的”——因为根本还没有任何作品、任何标准、任何可以被拿来比较的参照物。她不是在判断“这个句子好不好”，而是在征询“这里有没有东西”——而“有没有东西”的标准，在那一刻，除了她自己的身体感受之外，不存在于世界上的任何其他地方。一个物理学家的“这里有东西”，最终将被迫在实验数据面前证明或证伪自己；而C的“这里有东西”，它的证明恰恰就是它所能推进出来的那个作品本身——在此之前，没有任何外部标准可以为它背书。

第二，**反馈的来源与校准的闭环结构**。专家直觉的反馈来自外部世界——棋局输了、音错了、实验数据不支持。这种外部反馈使得专家的内部判断可以被校正：一个错误的直觉会在外部反馈中被标记为错误，从而在下次被修正。内部校准的反馈却来自创作者自己的身体，而且它构成了一个**闭环**——C判断一个片段“留还是不留”的依据，是“身上不对劲”或“好像离那个东西近了一点”，但判断这个“不对劲”本身是否“对”的标准，又在哪里？同样还是她的身体。这不是一个可以通过外部验证来打破的循环，而是一个必须在内部被持续推进的过程。创作者无法跳出自己的身体，去确认那一刻的“不对劲”究竟是真正的校准信号还是噪音——她只能继续试探，在下次、下下次的身体反馈中，让那个信号自己露出它的性质。专家直觉是一个可以被外部打破的校准循环；内部校准是一个只能在内部被推进的生成循环。前者的理想是“对应”——内部判断与外部标准对应；后者的逻辑是“凝定”——内部信号在反复试探中被推进为一个清晰的结构，这个结构本身成为了事后的“标准”。

第三，**功能的指向与被校准对象的本体论地位**。这一层是内部校准不可还原性的核心所在。专家直觉处理的对象，无论多么新颖——一个从未见过的棋局、一个从未听过的异常音、一组从未被观测到的异常数据——在本体论上都是**已经被给定的**。它们

是一个在公共世界中已经被摆在专家面前的东西。专家直觉的功能，是对这个已经被给定的东西进行**识别与判断**——“这个局面意味着什么”、“这个音偏了多少”、“这组数据里藏著什么”。内部校准处理的对象，则恰恰是**尚未被给定的**。C身体中的那个刺痛，不是一个已经摆在她面前的“对象”，而是一种在她内部的、尚未获得任何外部形式的张力。它不是一个待解决的问题，而是一个待生成的问题。内部校准的功能，不是去识别它“是什么”，而是通过反复试探的动作，**把它推进为一个可以成为外部对象的东西**——一个可以被陈述的语义结构、一个可以被写下来的小说。这就是为什么内部校准不能被还原为“前反思的具身判断”或“高阶专家直觉”——因为后两者处理的，终究是一个已经在世界上占据了某种位置的对象（哪怕这个位置是刚刚被数据揭示的）；而内部校准处理的，是一个**尚未从主体内部过渡到世界之中的张力**。它的核心动作不是“判断”，而是**推进**：把一个还不是东西的东西，做成一个东西。

那位物理学家所经验的“这里有东西”的直觉，当他凝神追问时，他追问的对象——那组异常数据——已经在他的意识之外、在实验仪器的屏幕上存在著了。他的直觉是对那个对象的反应，尽管那个对象的意义尚未被解读。而C的刺痛，在它浮现的那一刻，除了在她自己的身体之中，不存在于世界上的任何其他地方。它没有形状，没有名字，没有可以被另一个人观测到的物理踪迹。把这个刺痛从一种私人的、内部的、前对象的张力，推进为一个可以在公共世界中被阅读、被讨论、被批评的文本——这个动作，是所有专家直觉的运作中都找不到对应物的，因为专家直觉的起点处，对象已经在那里了，而内部校准的起点处，除了创作者自己的身体感受之外，一无所有。

基于这三重差异，本文对现有文献的理论贡献可以精确界定为：内部校准这一概念所补充的，不是刻意练习文献已经覆盖的“熟练判断”，甚至不是它的高阶形态“在复杂情境中的快速模式识别”；它所指认的，是发生在所有领域性判断之前的一个更原初的环节——**在对象尚未被给定时，身体感受如何通过反复的自我试探，把一个私人的、前对象的张力推进为一个可以进入公共世界的作品**。这一环节，是现有专家直觉文献未曾、也无法覆盖的空白地带。

（三）内部校准与既有邻近传统的对话：深感、生成-探索模型

在进入创造力认知研究之前，必须先向一个更贴身的传统交代：尤金·简德林（Eugene Gendlin）的“深感”（felt sense）与“聚焦”（focusing）。简德林在《体验过程与意义的创造》（1962）和《聚焦》（1978）中描述了一种前概念的、关于某个处境的整体身体感——它先于语言，却能对试探性的言说给出“不是这个／就是它”的身体裁决；当一个说法“对上”时，身体会发生一次可感的“松动转换”（felt shift），意义由此被“向前携带”（carrying forward）进符号。任何熟悉这一传统的读者都会立刻看出：本文第三节链条的第一环与第四环——身体标记的浮现、依身体反应定去留的反复试探——与简德林的描述高度共振。本文不声称首次发现“身体先于语言知道‘不是这个’”；这一发现属于简德林。

本文的增量在别处，而且恰好是简德林的框架无法安放的两处。其一，聚焦是一种**环境中性的个体技法**：治疗室、独处、一位不给答案的倾听者——它的全部设定里没有对手。本文解剖的却是这套身体技法在一个特定技术条件下的命运：当一个亚秒级的外部系统在 felt shift 尚未发生之前，就反复递上“已经对了”的版本时，聚焦所依赖的那个“停留在说不出之中”的时间窗口被结构性占据了。简德林描述了链条如何运转，本文解剖的是链条的**发生时机如何被占位**，以及占位之后长出的新回路——替代性校准与“对平滑的刺痛”。其二，聚焦作为一套可教授的自助步骤，预设主体可以随时召唤这一过程；而第六节的递归困境恰恰表明，“去聚焦吧”这一建议在替代性校准被反复强化之后会落空——能听见“该聚焦了”的那个身体感受器，正是被覆盖的对象。把简德林的技法直接开成 AI 时代的药方，等于要求病人用已被置换的器官去执行治疗。

因此，本文与简德林传统的关系是：借它确认内部校准链条在个体经验层面的实在性，同时越出它，去处理它从未面对过的问题——校准的技术条件，以及校准方式本身的可替代性。在此基础上，再来看创造力的认知研究传统。

在上述与专家直觉的对话之外，内部校准与认知科学创造力研究中的“问题生成”文献之间，也需要做出清晰的定位。

长期以来，创造力的认知研究——从Guilford（1967）的发散思维经典模型到Finke等人（1992）的生成-探索模型——集中在“给定一个问题，如何产生多样化的、新颖的解答”这一框架内。酝酿效应研究的是“暂时放下问题后解决率提升”的机制；发散思维研究的是“从一个提示词产出尽可能多的不同回答”的流畅性与变通性。这些研究的共同前提是：问题已经被给定。在一个典型的创造力实验中，受试者被要求“想出一个砖块的尽可能多的用途”——砖块和“用途”这两个词，已经把问题的边界划好了。

内部校准链条的前三个环节——身体标记浮现、松散吸附、跨域试探性联结——全部发生在“问题”被清晰陈述之前。C在最初感受到刺痛时，她面对的不是“如何用一个新颖的方式写窗户里的灯光”这个可以被清晰陈述的任务。她面对的是一种说不清楚的、尚未获得任何问题形式的张力——它甚至还没有成为一个“关于什么”的追问。内部校准的核心功能，恰恰是把这个还不

是问题的身体感受，推进为一个可以被追问的清晰问题（“距离”）。这是一个“问题生成”的过程，而非“问题解决”的过程。

Finke等人（1992）[3]的生成-探索模型已经触及了前语义结构（“前发明形式”）向显性概念的转化，这一点值得肯定，它也确实是现有创造力文献中与内部校准链条最接近的理论框架。但这里有一个关键的区分需要做出：Finke模型中所谓的“生成”，发生在一个已经被给定的概念域之内——受试者被要求“组合这些形状来创造一个有用的物体”，“有用的物体”这个框架已经在那里了。而在C的案例中，最初的身体标记连“它将成为小说的一个主题”这层归属都尚未确定——它可能成为小说，也可能成为一首诗的核心意象，也可能只是在日记里被记下一句话之后就消散了。内部校准的发生逻辑，比生成-探索模型的“在给定的域内生成变体”要更靠前一步：它追问的是，那个将要成为生成域的“域”本身，是如何在身体层面被第一次标记出来的。更关键的是，内部校准所处理的材料——那种前语言的身体张力——与生成-探索模型所处理的材料（形状、图形等可以在外部被给定的元素），在转化逻辑上有根本不同。生成-探索模型的“生成”，是受试者**主动地**在给定的材料之间尝试组合；内部校准的“推进”，是创作者**被动地**感受身体的微小反应，并在这种感受中让材料自己逐步露出它的结构。前者的逻辑是“探索与发现”，后者的逻辑是“试探与凝定”。这一步，是现有创造力研究尚未系统打开的地带。

综合上述两组对话——与专家直觉文献的差异论证、与创造力理论的问题生成定位——本文确立“内部校准”这一命名的不可还原性。它不是对既有概念的替代或包装，而是对既有文献共同留下的一块空白地的开垦：在公共对象尚未被给定、公共标准尚未存在、解答空间本身有待生成的环节，身体感受如何成为推进新结构发生的引擎。

四、替代：AI如何催生一种新的校准方式

第三节从正面建构了一条完整的内部校准链条。但要论证AI的介入催生了一种替代性校准方式的发生，仅指出“AI提供即时平滑替代方案”是不够的——这个判断停留在定位层面，尚未进入显微镜级别。本节将做一次更精细的追问：当C在那一刻点击“润色”按钮时，具体发生了什么？在那几秒钟内，C身体感受的质地经历了怎样的变化？这个变化如何从一次的“被平滑”，在反复累积之后，发生为一种新的内部姿态？

（一）一次替代的逐环解剖

回到C的案例。假设C在写下几个试探性片段之后，将其中一个粗糙的句子——“灯光穿过窗玻璃的时候好像被什么东西弯曲了，不是物理上的弯曲，是另一种”——贴进了AI辅助写作工具，点击了“润色”按钮。AI在不到一秒的时间内返回了几个版本。其中一个版本是这样的：“灯光穿过窗玻璃时，仿佛被某种不可见的力量弯曲了——不是光学意义上的折射，而是另一种更深层的扭曲。”

这个版本是好的。语法流畅，修辞得体，“不可见的力量”和“更深层的扭曲”甚至为原句增添了某种文学性的分量。C读了，觉得“不错”。但恰在此时——这是最关键的时刻——C的身体深处有什么东西动了一下。不是清晰的反对，不是明确的“这不对”。而是一种类似于“一根极细的弦被轻轻拨了一下随即被按住”的感觉。那个粗糙的原句中“不是物理上的弯曲，是另一种”——这个突如其来的中断、这个“另一种”后面没有被填上的词——在C的身体里原本制造了一种微小的张力，一种向外试探的冲动：那里有一个空缺，需要被填上，但填上什么，还没有人知道。这个空缺本身，就是C在那个当下最真实的状态的肉身化——她知道有什么东西在那里，但她说不出那是什么。那个句子的“未完成性”，正是对“说不出”这个状态的忠实记录。

现在，AI的版本用“更深层的扭曲”填上了这个空缺。C读到“更深层的扭曲”这五个字的时候，她的身体里发生了一件事。那根被拨动的弦，没有被弹响，而是被轻轻地按住了。那个原本向外试探的、对著一个空缺摸索的冲动，失去了它的对象——空缺已经不在。在那个空缺的位置上，现在站著一个完整的、流畅的、甚至略带文学光泽的短语。C的身体感受在那一秒内经历了一次质地的转变：从一种向外试探的、带有不确定的微小兴奋的刺痛，转变为一种向内的、轻微的、几乎不可察觉的塌缩感。不是剧烈的失望，不是清楚的“这不对”，而是——一种说不清楚的、好像有什么东西没有被碰到、但又说不出具体是什么的闷。最微妙的是，这种闷感在产生的同时，就已经被另一种更强烈的感受所覆盖：那个AI版本是好的。从任何可被公共评估的标准来看——语法的规范、修辞的得体、逻辑的连贯——它都是好的。C的身体里同时有两个信号：一个是那根弦被按住之后残留的微弱的闷，另一个是面对一个“好版本”时的轻微愉悦——省去了麻烦的愉悦。这两个信号的方向是相反的。一个说“有什么不对”，另一个说“这挺好的”。而这两个信号在力量上是不对等的：那个“不对”需要时间才能被推进成一个清晰的判断——C需要重新回到自己的粗糙句子，反复读它，让身体再感受几次，才能说出“更深层的扭曲”究竟错在哪里；而那个“挺好的”是即时的、直观的、有整个AI界面和社会常识做背书的。在这种不对等的张力中，C的意识做出了一个微小的决定——她保留了AI的版本。这个决定不是被迫的，不是被强制的，而是在她自己身体的两个信号之间，那个更清晰、更即时、更“合理”的

信号压过了那个更模糊、更缓慢、更无法为自己辩护的信号。

但这个决定做出了之后，C的身体里还发生了一件更隐蔽的事。那根被按住的弦，没有就此消失。它留下了一个极其细微的残余，一种“有什么东西没有被碰到”的未完成的感觉。然而，在接下去的创作中，C没有再回到这个残余。它被搁置了。下一次，下下一次，当类似的情境——写下一个粗糙句子、点击“润色”、看到一个流畅版本——再次发生时，那根弦被按住的过程变得越来越快，越来越不被注意。那个从“向外试探的刺痛”到“向内塌缩的闷”的质地转变，在一次又一次的重复中，不再被感受为一个事件，而是变成了创作流程中的一个常规步骤——就像你每天都要按下的那个电梯按钮，你不会再去感受它的触感。随之而来的，是另一条神经通路的逐渐强化：身体标记浮现→打开AI→看到平滑版本→接受。这条通路不是“什么都不做”的空白，而是一种新的内部动作：它要求创作者在那一瞬间，**调整自己的感受，使之与那个平滑版本相匹配**。那个最初C在重读自己粗糙句子时需要动用的能力——去感受“身上不对劲”或“离那个东西近了一点”——正在被一种不同的内部动作所替代：看到AI版本，判断它是否“合理”，将其整合进自己的文本，然后继续前进。这不是“校准能力”的废用，而是**校准方式的替代**——从一个向身体深处征询方向的内向校准，转变为一个对外部平滑选项进行快速评估与接纳的外向校准。

（二）替代性校准的三个可观察特征

基于上述解剖，本文将这种在AI即时平滑化介入下被反复强化出来的新的校准方式，命名为**替代性校准**。它的发生不是一次性的“截断”，而是一种新的感知-行动回路的逐步强化。它具有以下三个可被观察的特征。

第一，**校准方向的外向化**。在内部校准中，校准的方向是内向的——创作者向自己的身体感受征询方向，身体感受是校准的最终依据。在替代性校准中，校准的方向发生了逆转：创作者首先面对的是一个已经被外部给定的选项（AI的平滑版本），他的工作不再是从身体中“凝定”出一个语义结构，而是评估这个外部选项是否“可以接受”——是否符合语法、逻辑、可读性等公共标准。身体感受仍然参与其中（“我感觉这个版本不错”或“我觉得有点不对”），但它不再是最初的、唯一的校准源。它降格为对一个外部产物的二级评估。

第二，**校准速度的脱节**。内部校准的节奏由身体感受的浮现与推进速度决定，这个速度是缓慢的、反复的、无法被外部标准所催促的——C从感受到刺痛到最终凝定出“距离”，用了数天到一周多的时间。AI的替代性校准是亚秒级的：点击按钮的瞬间，一个完整的、平滑的版本就出现了。这两种速度之间的落差，并不仅仅是量的差距。它造成了一种结构性的不对等：当外部选项已经出现时，内部校准甚至还没来得及开始——C的身体需要更多次的试探，才能从那个粗糙句子中“读出”它究竟“身上有什么”，而AI已经给出了三个可以立刻被评估的版本。在这种不对等中，替代性校准的时间优势不是被“选择”的，而是作为一种既成事实，在内部校准启动之前就已经占据了那个空缺。

第三，**校准信号的质变**。内部校准所依赖的核心身体信号，是一种向外试探的、带有不确定的微小兴奋的刺痛——它在身体感受与未成形的语义结构之间保持著一种“被拉向某处”的张力。替代性校准在反复发生之后，创作者的身体里开始出现一种新的信号：一种面对AI平滑版本时的、轻微的、难以言说的抵触——不是对某个特定版本的抵触，而是对“总是有这么一个版本在那里”这件事本身的抵触。本文将这种新的身体信号称为**对平滑的刺痛**——它不是来自某一个空缺的初次浮现，而是来自那个空缺总是在被填上之前就被平滑掉的累积性经验。这个对平滑的刺痛，是替代性校准发生过程中同时在创作者身体里催生出来的一种新的敏感带。它说明，替代性校准不是一个单向的“堕落”过程，而是一个充满内部张力的发生场——旧的校准方式被覆盖，新的抵触在生成。这个新信号，正是下一轮生成循环的可能起点。

五、替代性校准发生的三重后果

如果第四节的论证成立——即AI的即时平滑化正在催生一种替代性校准方式的发生——那么接下来的问题就是：这种转变在实际的创作世界中会产生什么样的后果？本节从创作者个体、作品品质和创作生态三个层面，逐层展开论证。

（一）创作者个体层面：一种新的内部张力状态的生成

对个体创作者而言，替代性校准的发生并未导致“写不出东西”——恰恰相反，借助AI工具，创作者的产出频率和文本数量可能显著提升。更准确的描述是：创作者正在经历一种新的张力状态。他的身体仍然偶尔浮现那些微小的标记，冬夜的灯光仍然会在某些时刻刺痛他。但现在，在那些标记被AI平滑化的反复经验中，他身体里同时生成了两种互相矛盾的信号。一种是旧有的、向外试探的刺痛——它告诉他“那里有东西”。另一种是新的、对平滑的抵触——它告诉他“那个东西总是被平滑掉”。这两种信号同时存在、方向相反，构成一种独特的内部摩擦：创作者既被刺痛的引力所吸引，又被对平滑的抵触所阻挡，在两种力量之间疲惫地来回。这不是“灵感枯竭”——枯竭意味著不再有东西涌现，而这里是东西在涌现，但接住它的那条链条已经被另

一条更快的链条所覆盖。这种状态比枯竭更消耗人，因为它不是信号的缺失，而是信号与信号之间的冲突。

（二）作品层面：推进痕迹的稀释

在作品层面，替代性校准的发生导致一种质地的转变。每一篇作品，除了它所传达的语义内容之外，还携带著一种可以被感知的“身体质地”：语言在推进过程中留下的那些微小的不对称、那种节奏上无法被完全正当化的停顿与加速、那些不是为了传达信息而是仿佛作者身体在文字间呼吸的痕迹。这种质地不是“风格”，因为风格可以被刻意经营甚至被模仿；这种质地是那条内部校准链条在作品中留下的“推进痕迹”——就像一块手工玻璃中那些不规则的微小气泡，它们不是缺陷，而是制作过程中玻璃与工匠的身体之间实际发生过互动的证明。

当替代性校准在微观层面反复发生时，被改变的不是作品是否“流畅”，而是作品获得那种推进痕迹的发生路径本身。一条从身体感受内部被凝定出来的句子，它的节奏、它的词语选择、它那些无法定义为“偏差”的微小特质，是那条被走完的推进链在语言表面留下的涟漪。而一条经由替代性校准整合进文本的AI句子——即使被创作者仔细修改过——却不是从那条链条中被推进出来的。它来自一个外部生成、内部评估的过程，它的质地来自AI的训练数据与创作者的筛选判断之间的耦合，而不是创作者身体感受的反复凝定。创作者仍然可以产出高品质的作品，但那些作品获得身体质地的路径，正在从“内部推进”转向“外部筛选与整合”。这不是“机器写的”和“人写的”之间的二元差异。一个大量使用AI辅助写作的创作者，他最终定稿的文本可能经过反复修改、注入了大量个人判断，但它获得身体质地的那个最原初的身体起点——那个最初粗糙的、在语言边缘颤动的、尚未决定自己将成为什么的状态——已经在进入文档的最初几秒，就被一条更平滑的路径所覆盖了。

（三）生态层面：练习场的功能转变

最深层的后果发生在创作的生态层面。一个健康的创作生态，不仅需要作品被生产和消费的循环，更需要无数不被看见的、无关紧要的、与产出没有直接关系的创作活动——那些写了又删掉的碎片、那些漫无目的的练习、那些仅仅为了维持手感而进行的反复操练。这些活动构成了内部校准能力的练习场。一个写作者的内部校准能力，不是在那些高光的、决定了作品成败的关键时刻被训练出来的，而是在这些看不见的、不被计入“产出”的日常操练中，被一点一点养起来的。

当AI工具将“省去不必要的麻烦”内化为创作流程的默认设置时，最先被省去的恰恰是这些看不见的练习。因为它们本来就没有明确的目的，本来就不在效率计算的视野之内。一个写作者在过去可能花一个下午写一些明知不会成文的东西，只为保持手感；如今，这个下午因为可以“用AI快速搞定正事”而被解放出来——但被解放出来的时间，并不会被重新投入到那些无目的的练习中。不是因为创作者“变懒了”，而是因为替代性校准被反复强化的过程中，那些练习的价值本身变得不可见了。它们从“维持手感的必要功课”被重新定义为“效率低下的时间浪费”——这个重新定义不是任何人有意为之的，而是替代性校准在发生过程中，同时重塑了创作者对“什么是有价值创作活动”的感知本身。

练习场萎缩的后果，不是即时的。它不会让下个月的文学期刊变得难看，不会让今年的畅销书榜单缺少佳作。它的效应是滞后的、生态性的：当一代新创作者不再经历那些不被看见的、无目的的、仅仅为了维持手感而进行的练习之后，他们作品中能被感知到的那种“从身体内部被推进出来的质地”，会变得越来越陌生——不是因为他们不想要它，而是因为他们已经不知道那种质地是如何被获得的。他们学到的创作方式是：从阅读已有作品和自己的AI辅助实践中，学会判断什么是“好作品”。但成品可以被模仿，成品背后的推进链无法从成品中学到。生态系统失去了一条从地下蓄水层汲取水源的隐秘通道，地面的河流还在流淌，但水质在发生无法被一眼看出的变化。

六、反驳与回应

上文描述了替代性校准的发生及其后果。这一论证在展开过程中，必须面对两个最具威胁性的反驳。本节正面回应这两个挑战。

（一）反驳一：“流畅工具无罪论”——更好的工具不该为使用者的选择负责

这个反驳可以这样表述：本文把替代性校准的发生归咎于AI工具，但这是否混淆了“工具提供了便利”与“使用者放弃了练习”？计算器没有让数学家失去数学直觉，因为数学家知道什么时候该放下计算器、回到纸笔推演。一个严肃的创作者同样可以选择在某些时候使用AI的润色功能，在其他时候关掉它，进行不被打断的独立试探。如果创作者因为工具的便利而让自己的内部校准被替代，那是创作者自己的选择问题，不是工具的问题。一个成熟的创作者，完全可以在使用AI的同时，刻意地、分阶段地进行不被AI干扰的练习——就像运动员既使用高科技设备，也进行无装备的基础训练一样。

本文对此回应如下。

首先，承认这个反驳的部分有效性。确实存在创作者能够自律地使用AI工具、不让它覆盖自己的核心校准过程。本文从未主张AI工具是唯一的、不可抵抗的决定因。本文的主张是：AI工具所创造的“平滑”体验，因其**即时性、愉悦性和界面暗示**，在结构上比计算器或打字机这类工具更倾向于诱发替代性校准的发生。计算器提供的是结果——一个精确的数字，使用者仍然需要自己去理解和应用这个数字；AI“润色”提供的是替代品——一个完成了的、看起来更“好”的表达，这种替代发生在创作者对自己原句的身体感受尚未成形之前。两者对使用者认知过程的介入深度不同。

但这不仅是程度问题。这里存在一个更关键的机制差异，需要被明确指出。反驳者说：创作者可以“刻意地、分阶段地”进行不被AI干扰的练习，就像运动员既使用高科技设备也进行无装备基础训练。这个类比忽略了一个根本性的不对称。运动员的“无装备训练”与“高科技训练”之间，不存在一项技术在**亚秒级的时间尺度上**占据另一项训练的启动位置的问题。运动员可以在上午做无装备训练、下午用高科技设备，两者在时间上被清晰地分隔。但创作过程中的内部校准——特别是第四环的身体校准反复试探——不是一个可以被安排在“某个特定时段”进行的独立活动。它恰恰发生在那些未被规划、未被预定的瞬间：一个粗糙句子刚被写下的那一刻，身体感受刚刚开始向那个句子试探的那一刻。AI的即时平滑化介入的，正是这个瞬间——它在内部校准尚未启动之前，就给出了一个替代方案。创作者无法“选择”在这一刻不使用AI，因为在这一刻，他的身体校准甚至还没有来得及告诉他“这里需要一次练习”。AI占据的，不是练习的“时间”，而是练习的**发生时机**——那个在身体感受刚要向外试探、尚未找到方向的最初瞬间。这个时机一旦被反复占据，替代性校准就在创作者意识到“练习”的价值之前，一步一步地被强化了。

这就触及了反驳者自身的一个隐含预设。这个反驳预设了创作者是一个自主的、理性的、能够清晰区分工具使用与能力练习的主体。但第三、第四节的案例解剖已经表明，创作过程中最关键的校准发生在一个前反思的层面——那个身体感受刚刚浮现、尚未被意识清晰捕捉的边缘地带。AI工具的平滑介入，恰恰作用在这一层面：它在创作者尚未完成身体校准之前，就已经给出了一个替代方案。创作者不是做了一个“偷懒”的决定，而是在决定被意识到之前，就已经被引向了那条更平滑的路径。这不是意志力的失败，而是时间结构的不对等——替代性校准的速度远远快于内部校准，以至于后者总是在启动之前就失去了它的对象。

最后，这里存在一个递归困境，需要被明确指出。反驳者说：创作者可以“自律地”选择不使用AI润色功能。但“自律”——作为一种在外部诱导面前保持自身判断方向的反思能力——本身就是内部校准链条所维持的那种身体敏感性的认知延伸。换言之，一个创作者能否在“AI改得挺好”的即时诱惑面前，依然有能力听到自己身体深处那一声微弱的“不对劲”，并根据这一声“不对劲”采取行动——这恰恰取决于他的内部校准能力是否还足够敏锐。如果这条链条已经在之前数十次、数百次替代性校准的发生中被覆盖了，那么那种能够触发“自律”的身体信号，就可能在创作者意识到之前便已消散。在这种情况下，要求创作者“自律”，就像要求一个已经被训练成另一种校准方式的创作者，去使用他已经不再拥有的那种校准方式来做出选择。不是他不愿意，而是他用以做选择的那个身体感受器，已经在反复被覆盖之后，变成了另一种东西。

（二）反驳二：“高阶专家直觉可覆盖论”——内部校准不过是专家直觉的创作版

第二个反驳更危险，因为它直接动摇了本文核心命名的不可还原性。反对者可以这样表述：本文费力论证了内部校准与专家直觉的差异，但这个差异真的站得住脚吗？一个顶尖科学家在面对全新异常数据时的“这里有东西”的直觉——那个数据从来没有人见过，它的意义完全未知，在他的领域内没有任何现成的标准可以评判它——这种直觉，与C的“刺痛”，在本体论上究竟有何不同？两者都是在标准空缺处向自身长期训练所积淀的身体化感知征询方向。如果本文只是把“专家直觉”这个概念在创作领域重新描述了一遍，那“内部校准”就只是一个多余的标签。

本文在第三节已经从三个层面论证了内部校准与专家直觉的差异（标准的性质、反馈的来源与校准的闭环结构、被校准对象的本体论地位）。在此，面对这个更尖锐的版本，本文需要将论证再推进一层。

那位科学家所面对的“异常数据”，无论多么新颖、多么未被解读，在本体论上有一个决定性的特征：它已经是一个**公共对象**。它存在于实验仪器的屏幕上，可以被另一个科学家观察到，可以被复制、被讨论、被反驳。当科学家的直觉告诉他“这里有东西”时，他的直觉所指的对象，是一个已经从自然世界中进入到人类公共知觉领域的物理信号。他的直觉是对那个公共对象的反应——尽管他尚未理解那个对象意味着什么。他接下来的工作，是用各种理论工具和方法去**解读**那个公共对象，去揭示它隐藏的意义。他的直觉的成败，最终可以被那个公共对象的后续行为所验证。

C的刺痛，则在本体论上处于一个完全不同的位置。它是一个**纯粹私人的、尚未获得任何外部形式的内部事件**。它没有形状，没有名字，没有可以被另一个人观测到的物理踪迹。它不仅“尚未被解读”，它甚至还不是一个可以被解读的“对象”——

它是一团混沌的、没有边界的、附著在C的身体上的张力。C接下来的工作，不是去解读一个已经在那里的对象，而是**把这个还不是对象的内部张力，推进为一个可以进入公共世界的对象**——一个可以被写下、被阅读、被讨论的文本。在这个推进完成之前，那个刺痛不属于公共世界；它只属于C，而且是以一种连C自己也无法清晰说出的方式属于她。

这就是内部校准与高阶专家直觉之间不可跨越的本体论鸿沟。专家直觉始终是对**一个已经在公共世界中占据了某种位置的对象**的反应性判断——无论那个对象多么新颖、多么未被理解。内部校准则是**将一个尚未从主体内部过渡到公共世界之中的张力，推进为一个可以被公共识别的对象的生成性过程**。前者的核心动作是“识别”（哪怕是创造性的识别——“我看出了这里有一个前人没有看出来的结构”）；后者的核心动作是“推进”（“我把我身体里的一个还不是东西的东西，做成一个东西”）。这两种动作的差异，不是领域的差异（科学vs艺术），而是**人与其对象之间关系的范畴性差异**。一个是对象先于判断，一个是判断的过程本身就是对象的生成过程。

如果“内部校准”可以被“高阶专家直觉”覆盖，那就意味著，在所有那些对象尚未被给定的时刻——在一个科学家还没有任何异常数据可以看、一个作曲家还没有任何音符可以听、一个小说家还没有任何意象可以想的时候——在那些一无所有的起点处，就已经有某种可以被称为“直觉”的东西，正在把一个尚未存在的东西，从虚无中召唤到存在。这不是“专家直觉”这个概念所描述的东西，这甚至不是任何以“判断”为核心动作的概念所能覆盖的。这就是为什么内部校准不能被还原——因为它所指的，是所有“判断”发生之前的那个更原初的环节：**当判断的对象本身还有待被生成时，创造者的身体在做什么。**

七、推衍：替代性校准发生的普遍性

本文以上的论证以生成式AI为直接的考察对象，但所触及的问题并不限于AI。如果替代性校准的发生机制如第四节所述——即当一个外部系统能够在创作过程的微观层面、在内部校准尚未完成之前就提供即时的平滑替代方案时，一种从内向校准向外向校准的替代便会在反复发生中被强化——那么在任何以“高效产出”为默认规范的创造性场域中，都可能观察到同类现象。本节在教育与学术两个领域进行简要推衍，以显示本文论证更广泛的辐射范围。

（一）教育领域：红笔的平滑逻辑

教育领域的替代性校准发生，在时间上远早于AI的出现。理解这一现象的关键，不在于把“标准化教育体系”当作一个现成的批判对象，而在于追溯一个微观的教学动作如何在特定的历史条件下被生成出来，并在无意中促成了一种类似于AI润色的替代性校准的强化。

设想一个二年级的小学生S，他正在写一篇与任何考试无关的随笔。他写下一个句子：“我家的猫走路的样子好像——”然后他停下了。他知道有一种很具体的感觉，但不知道用什么词。他试了几个：“很骄傲”？不对，猫走路不是骄傲。“很轻”？太普通了，不是他要的那个感觉。他划掉，重写，再次划掉。这个过程在一个旁观者看来“没有效率”——一个小孩对著一个没写完的句子发呆、反复涂改，显然比抄写一个标准答案浪费时间。但这个过程恰恰是内部校准的第四环在一个八岁孩子身上的运作：他正在用自己的身体感受反复试探那个空缺的形状，并在每一次试探中把那个“说不出的感觉”推进得更清晰一点。[4]

如果此时，他的语文老师走到他身边，看到他卡住了，于是弯腰在作文纸上用红笔划掉那个破折号，在旁边写下示范：“猫走路很优雅。”老师的意图是帮助，动作是微小的，效果是即时的——S的句子完成了，看起来“好”了。但被覆盖的不仅仅是一个不成功的句子。被覆盖的是S在那个当下正在进行的一次内部校准练习。他本该在“说不出”的状态中多停留一会儿，让身体感受把那团模糊的东西推到更清晰的地方——但老师的红笔在他内部校准尚未完成之前，就给出了一个平滑的替代方案。

这个场景中的“红笔”，不是一个需要被批判的邪恶工具。它的出现，本身就是近代教育大众化与效率化潮流中的一个历史产物——当一个教师需要同时面对四十个学生的作文时，她无法逐一蹲下来陪每个孩子慢慢试探那个“说不出”的东西。红笔的“即时平滑化”逻辑，正是从这种结构性条件中被生成出来的：它是一种在资源约束下，为了让写作教学能够被规模化执行而被发明出来的效率工具。它与AI润色功能共享同一个生成根源——都是用一个外部的、平滑的、即时的替代方案，来在内部校准尚未完成之前就给出答案。差别在于：红笔对内部校准第四环的覆盖，发生频率受限于教师的时间——一学期可能只有几次；而AI润色对第四环的覆盖，可以发生在每一次写作的每一句话上。红笔打开了替代性校准的路径，AI把这条路径变成了创作流程的默认设置。

到了中学或大学阶段，大量学生表现出一种典型的张力状态：他们能写出符合标准的、没有错误的文章——语法正确、结构清晰、论点合理——但当他们被要求写一些“真正属于自己的东西”时，一种难以言说的阻滞出现了。身体里隐约有东西，但通往语言的那条路找不到。这与第五节所描述的创作者的张力状态是同构的：身体标记仍然偶尔浮现，但那条把它推进为清晰语义

结构的链条，在多年被外部标准反复替代之后，已经被一种更高效的、对外部标准做出快速响应的校准方式所覆盖。

（二）学术研究领域：“这里有东西”的感受器如何被指标替代

学术研究中的替代性校准发生，有其特殊的表现形式。一个研究者在选择研究方向时，内部校准的功能是：在大量可能的课题中，那些让研究者身体感受到“这里有东西”的问题，会被优先选中。这个“这里有东西”的感觉无法被清晰论证——如果它已经可以被清晰论证，那它就不再是一个研究问题，而是一个已有答案的待办事项了。内部校准在这里的作用，与在艺术创作中如出一辙：它让研究者被某个尚未成形却隐隐刺痛的疑惑所吸引，投入长时段的、不确定能否产出成果的探索。

当前的学术评价体系——项目制、量化考核、出版指标——在宏观层面对这一过程构成了明显的挤压。一个可以稳定产出中上水平论文的方向，远比一个风险极高、可能若干年无产出的新方向更符合生存理性。申请基金时，研究者需要明确写出研究目标、步骤和预期成果——这对于需要在混沌中摸索的崭新方向而言，构成了一种结构性障碍：它要求研究者在出发前就描述他们只能在抵达后才可能理解的事情。

但更隐蔽的问题发生在微观层面。当一个研究者反复经历“按照指标要求写立项申请→获得资助→产出符合预期的论文→获得续期资助”这条闭环时，一种替代性校准正在他的内部发生。他不再需要向自己的身体感受征询哪个问题值得追问——指标体系已经给出了明确的优先级排序。久而久之，那些最初吸引他进入学术领域的、说不清楚但身体知道“很重要”的疑惑，被搁置在“等有空的时候再想”的角落里。但内部校准与艺术创作中的情况一样，是一项需要反复练习的技艺。当“向自己的身体征询方向”这个动作长期不被执行时，另一个动作——向外部指标征询方向——则在反复强化。学术产出仍然丰沛——论文数量、引用次数、项目经费都在增长——但越来越少的研究让同行读到时会产生那种“这个人触碰到了某些他自己也还在摸索的东西”的身体共鸣。那是内部校准被替代性校准覆盖之后，在学术生态层面留下的痕迹。

这组推衍共享同一个结构：在任何创造性场域中，当一个外部系统能够在认知过程的微观层面、在内部校准尚未完成之前就提供即时的平滑替代方案时，一种从内向校准向外向校准的替代便会在反复发生中被强化。这不是对某个特定技术或制度的指控，而是一个可以被观察、可以被检验的机制命题：它追问的不是“能力是否衰减”，而是“何种新的感知-行动回路正在被生成出来”。

八、结论：在承认替代性发生的前提下，重新设计不适感

本文完成了一项基础性的工作：提出并建构了“内部校准”与“替代性校准”这一对命名，用以追踪AI时代创作过程中一个此前未被精确把握的转变。本文论证，生成式AI——以最小化预测误差为训练目标——在被广泛使用为创作工具时，不仅压缩了创作所需的冗余空间，更为关键的是，它在创作者进行内部校准的微观过程中、在内部校准尚未完成之前反复提供平滑的替代方案，使得一条从外部选项出发、经由快速评估而整合入文本的外向校准链条，在反复强化中逐步覆盖了那条从身体感受出发、经由反复试探而推进为语义结构的内向校准链条。这一转变的核心，不是能力的“衰减”——那是一种从预设的理想状态出发的缺损诊断——而是一种替代性校准方式的发生：创作者正在从一个向身体深处征询方向的人，被重塑为一个对外部平滑选项进行快速识别与接纳的人。由此被生成的，不是一个“更少”的创作主体，而是一个拥有不同感知-行动回路的创作主体——其身体标记仍然会浮现，但接住标记的链条已经发生了结构性的替代。

与此同时，本文严肃回应了两个强有力的反驳。对于“流畅工具无罪论”，本文指出了AI工具与计算器这类工具在介入创作过程时机上的性质差异——AI的即时平滑化占据的不是练习的“时间”，而是练习的“发生时机”，它在内部校准尚未启动之前就给出替代方案；并指出了一个递归困境：反驳者所诉求的“自律”能力，本身就是内部校准所维持的身体敏感性的一部分，它在同一过程中被覆盖。对于“高阶专家直觉可覆盖论”，本文在第三节和第六节中论证了内部校准与专家直觉之间的本体论鸿沟——前者是对一个尚未从主体内部过渡到公共世界之中的张力的生成性推进，后者是对一个已经在公共世界中占据了某种位置的对象的反应性判断。

本文的贡献不在于提出一套操作方案，而在于为理解AI时代创作过程的转变提供了一个更精确的机制框架。它将讨论从“AI会不会替代人类创造力”的抽象争论，引向一个可以具体观察和检验的发生机制：在创作者与AI工具之间每一次微观互动的那几秒内，创作者的身体感受经历了怎样的质地变化？替代性校准的反复发生，在多长的时间尺度上会导致创作者内部校准的内向链条被外向链条系统性地覆盖？这些问题可以被转化为实证研究的对象。

然而，本文必须在此面对一个来自自身论证内部的严肃追问。如果本文拒绝了“保护”的框架——拒绝将内部校准视为一个需要被隔离保护起来的脆弱本质——那么，这篇论文在实践上的指向究竟是什么？如果问题不在于“回归”某种前技术的创作方

式，也不在于“保护”某种正在被替代的能力，那么我们能做什么？

本文的回应是：真正的实践指向，不是保护旧结构免受外部侵蚀，而是在**承认新结构已经在发生的前提下，创造新的矛盾，使系统无法停滞于单一的替代方向**。如果当前的问题是AI的平滑化几乎构成了唯一的方向——它把创作过程中那些不可避免的“不适感”（说不出的刺痛、写了又删的反复挫败、在混沌中摸索的漫长等待）都标记为“可以被一键消除的Bug”——那么，一个可操作的介入策略，不在于呼吁创作者去“忍受”这些不适（这又落回了意志力的预设），而在于能否设计出新的技术、制度或实践，来**重新激活不适感在创作过程中的生成性功能**。

这可能指向一些具体的探索方向。在技术层面：能否设计一种AI写作工具，它的核心功能不是“润色”或“优化”，而是“制造不适”——针对创作者的文本，生成那些让创作者身体隐隐不舒服的问题，比如“这句话为什么是这个节奏？”、“这个词和你前面用的那个词之间，真的没有更深的联系吗？”——从而**强化**而非绕过创作者向自己身体感受征询方向的动作？在制度层面：能否在写作教育中有意识地设计一些“不能被外部标准评估的练习时间”——不是为了“保护”一种能力，而是为了让学习者长期暴露在“没有任何外部标准告诉你做得对不对”的经验中，从而在反复的内部校准练习中，让那条内向链条获得足以与外向链条形成真实竞争的强化？在创作者自我实践层面：能否有意识地进行一种“粗糙练习”——刻意不润色、刻意保留那些让自己身体不太舒服的句子，并在重读中追问那种不适的质地，从而维持身体感受作为校准源的活性？

这些探索的具体形态，超出了本文作为一篇以基础诊断为任务的论文所能详述的范围。但本文愿意为这些未来的工作提供一个地基：如果不能先精确地指认那个正在发生的转变是什么——不是“创造力的衰减”，不是“灵感的枯竭”，而是一条从身体到语言的内向推进链正在被一条从外部评估到内部整合的外向链条所系统性地替代——那么所有的应对努力，都会在最开始的环节打偏方向。

最后，本文愿意以一个朴素而诚实的证伪条件收束全文。上述一切论证都可能被一个事实所动摇：如果我们观察到，一个在其日常创作流程中大量使用AI即时平滑化辅助、其微观创作决策高度依赖替代性校准的创作者，在长时段内（例如十年以上），仍然能够持续产出具备第四节所描述的那种“身体质地”的作品——那种让人读到时感觉“这个句子不是在脑子里被构造出来再经由外部评估整合进去的，而是在身体里被反复试探之后凝定出来的”的质地——那么本文关于替代性校准发生的核心判断便宣告失效。[5]因为这将证明，内部校准的内向链条，不需要通过反复的、不被替代性校准打断的练习来维持；它可以在一条被外向链条反复覆盖的创作流程中，依然保持其独特的生成力。在那一天到来之前，我们或许应当对自己身体里那条古老而低效的推进链，抱持一种即使在效率法庭上找不到充分理由、却也不愿放弃的审慎关注。这种关注，不是对一个逝去的黄金时代的乡愁，而是对一种正在被覆盖的身体技艺的承认——承认它不是在“衰减”，而是在被另一种更快的、更平滑的校准方式所替代。而这种替代本身，作为一个仍在进行中的、充满内部张力的发生过程，恰恰是下一轮追问的起点。

注释

[1] 材料说明：本文第三节所描述的创作案例（“写作者C”）系作者基于长期创作观察与反思而构建的思想实验性质的例示。其中的人物已完全匿名化，具体情境与时间线索均经过合成、简化与典型化处理，并非真实个案的记录。该案例仅用于揭示内部校准链条的生成逻辑，不应被视为经验数据或实证材料。关于思想实验与机制论证之间的张力，作者承认这一局限：思想实验能够展示一种逻辑可能性，但无法替代对真实创作过程的经验追踪。未来研究可考虑以自我民族志或对创作者的深度访谈为基础，为本文的论证提供更厚实的经验基底。

[2] 对话范围界定：关于内部校准与“刻意练习/专家直觉”理论的详细区分，见本节第（二）小节。此处预先说明：内部校准这一命名专指“在公共对象尚未被给定时，身体感受如何通过反复的自我试探，把一个私人的、前对象的张力推进为一个可以进入公共世界的作品的生成性过程”；它与专家直觉（对一个已经在公共世界中占据某种位置的对象的反应性判断）在被校准对象的本体论地位、校准动作的性质上存在范畴性差异，不可相互还原。

[3] 与Finke模型的关系：Finke等人的生成-探索模型（Geneptore model）是现有创造力认知研究中与内部校准链条最为接近的理论框架。本文肯定其触及前语义结构向显性概念的转化，但指出该模型所处理的“生成”仍发生在已给定的概念域之内，且其核心动作是对给定材料的主动“探索与发现”；内部校准所处理的则是尚未被域化的前语言身体张力，其核心动作是被动的“试探与凝定”。这一区分构成本文理论贡献的核心之一。

[4] 教育案例说明：第七节中有关小学语文教学的例示，是对常见课堂情境的典型化重构，旨在揭示“红笔”的即时平滑化逻辑与AI润色功能的结构同构性。该例示中的人物与情境均为虚构，仅作理论推衍之用。

[5] 证伪条件的操作化补充：本文第8节给出的证伪条件，并不要求观察对象完全排除AI工具的使用。其可检验的核心在于：如果在长期（十年以上）且大量使用AI即时平滑化辅助的条件下——即创作者的微观创作决策高度依赖替代性校准——创作者仍能持续产出具备第四节所描述之“身体质地”的作品（即作品语言中可感知到从身体内部被反复试探之后凝定出来的推进痕迹，而非经由外部评估后整合进文本的平滑质地），则本文的“替代性校准发生”命题即被证伪。这一条件的设定，是为了将本文的机制假说推向可被经验研究检验的方向。“身体质地”的判定本身可进一步操作化为盲评区分度：邀请对作者与创作流程双盲的资深编辑或作家，对两组文本——长期重度依赖即时平滑辅助的创作者的产出，与低依赖对照组的产出——做“推进痕迹”判别；若判别正确率长期不显著高于随机水平，且两组作品在独立质量评估中无差异，则本文命题被证伪。

[6] “身体感受”的术语说明：本文所使用的“身体感受”一词，在广义上与现象学传统——尤其是梅洛-庞蒂关于“身体图式”与“肉身化的意识”的讨论——有深层共鸣，但本文不进入现象学的技术性框架。简言之，本文的“身体感受”指涉的是一种前语义的、指向性尚未被确定的“存在张力”——它不是纯粹的生理信号（如饥饿或疼痛），不是情绪（如悲伤或喜悦），也不是被压抑的无意识内容。它是一种在意识边缘涌现的、主体在场却尚无法清晰命名的冲动——一种“有什么东西在那里但我还不知道那是什么”的身体确信。这一界定区别于Damasio的“躯体标志”假说（后者专指决策过程中将情绪信号与预期结果联结起来的身体状态），也不同于认知科学中的“内隐加工”概念（后者指向的是在意识阈限之下自动运行的认知过程）。本文的身体感受是主体可感知的——虽然它尚未获得清晰的形式——但它不是无意识的，而是处在一个“即将被意识捕捉但尚未被捕捉”的过渡地带。正是这个过渡地带的存在，使得内部校准成为可能：创作者在反复试探中，将这个地带中的张力逐渐推进为清晰的语义结构。

参考文献

- Anderson, M. L. (2010). Neural reuse: A fundamental organizational principle of the brain. *Behavioral and Brain Sciences*, 33(4), 245–266.
- 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.
- Gendlin, Eugene T. (1962). *Experiencing and the Creation of Meaning*. New York: Free Press of Glencoe.
- Gendlin, Eugene T. (1978). *Focusing*. New York: Everest House.
- Finke, R. A., Ward, T. B., & Smith, S. M. (1992). *Creative cognition: Theory, research, and applications*. MIT Press.
- Guilford, J. P. (1967). *The nature of human intelligence*. McGraw-Hill.