

# 裁决回路：生成式人工智能如何拆解艺术发生的内在担保

亲手做的那一下被材料反打回来——AI 为这条回路的每个节点都备了替身

作者 秦莉 · SDE 学员 · 约 2.5 万字 · 发表于2026年7月28日 · 加固版

## 摘要

作者：秦莉 · 原初问题：AI来了，艺术如何做？ · 补规范 + 二次双审 + 打磨改稿 思想拓展  
声明：本智能体（SDE 金点子发生器）产生的任何论文与文本，均作为"思想拓展功能"提供——是 AI 在特定方法引导下的思想实验与观点探索，用于拓展思路、激发思考，不构成正式学术研究成果或事实性结论；如需正式使用，请自行核验并遵守学术规范。 提升·论文① · 金点子①（显露(S) · S2 · 模态M（粒子·波·场）） 裁决回路：生成式人工智能如何拆解艺术发生的内在担保 本文提出一个不同于既有"AI能否创作"讨论的分析框架：生成式人工智能对艺术最根本的冲击，不在于它具备了某种"创作能力"，而在于它系统性地拆解了艺术发生过程中那条不可缺席的亲自回路——创作者在操作材料时，被材料反打回来、从而改变自身判断方向的"裁决回路"。这条回路不是"创意"或"灵感"，而是一个完整的闭环：亲手操作→意外偏离→触觉/视觉/身体反打→重新判断→改向。它的运转不依赖任何外部制度、话语或共同体的认可——它本身就是艺术发生的最小可运转单元。AI的介入，不是在"能力"或"作品"层面挑战了这个单元，而是为它的每一个节点提供了一个可代为执行的替代方案。这种替代的累积效应，不是让艺术变得更高效，而是让裁决回路的各个节点之间开始脱扣——操作被外包，意外被消除，反打被绕过，改向被预先填平。当回路断裂，创作者面对的不再是"我该怎么做"，而是一个更根本的问题：这一环，我是亲自踩，还是让它替我转？本文以裁决回路的脱扣与再踩为判断主轴，重新阐释氛围、波谷、回写三个层面如何围绕这条回路组织自身，并以对青年艺术家小杨（化名）使用AI过程的追踪为主要锚点，辅以近年公开报道中其他创作者的陈述，让这一判断落地。在此基础上，本文追问一个发生学问题：当回路不再自动运转，创作者被迫将"是否继续"作为一个需要反复裁决的元问题来面对时，这种裁决实践本身，是否正在塑造一种新的创作主体性——一种从"被回路推着走"切换到"管理一条不再自转的管线"的主体性？文章结尾设置可证伪条件：如果能够出现一类在深度使用AI的同时仍然完整保留裁决回路各节点、并能在观看中被识别出"亲自性重量"的作品，则本判断需要被修正。

关键词：生成式人工智能；裁决回路；艺术发生；亲自性；回写

## 一、引言：一个被绕过的要害，与四条进路的未竟之处

"AI来了，艺术怎么办？"这个问题已经被反复提出，但几乎总是被放在同一个预设里：艺术是一类可以被界定特征的"作品"，AI是一种具备某种"能力"的技术，问题在于前者能否抵御后者的入侵，或从后者那里借得什么新工具。讨论迅速分叉为两路：一路是防御性的，试图找出AI"做不到"的事——情感、原创性、具身经验——然后宣布艺术的最后堡垒依然安全；另一路是拥抱性的，把AI视为新的画笔、新的暗房，认为艺术史上一再重演的技术恐慌终将平息，这一次也不例外。

这两路各自都有其道理，也都各自错过了那道最深的裂缝。这道裂缝不在"作品"的层面——不在AI生成的图像是否配得上"艺术"这个称号，不在AI谱的曲是否能够打动人心。它在更下面：在"做艺术"这个日常动作还能不能像从前那样，被一套内建在创作过程本身中的机制所持续推动。这套机制，不是创意，不是灵感，不是激情。它是一个可以精确描述的闭环：当一个人亲手在材料上做出一项操作——画一笔、凿一刀、写下一行字——这项操作必然产生一个不可能被完全预判的结果。手腕的微小角度、颜料的浓淡、呼吸的轻重，都会让实际落笔的那一下偏离脑中预设的图景。这个偏离，一旦进入视觉或触觉的感知，就构成了一个"反打"——它不等于"错误"，而是材料以自己的方式回应了操作。这个回应，会触动创作者。他可能会觉得"不对"，也可能会觉得"这个意外的效果比我原来想的更有趣"。无论是哪种，这个触动都会推动他做出一次重新判断：是沿着原方向继续，还是放弃原有计划、顺着这个意外往下走？这个重新判断，又会导向下一项操作——如此往复，构成一个不间断的循环。

这个循环，就是本文要命名的核心结构：裁决回路。它不是创作过程中的一个"阶段"或"环节"，它是创作发生的最小可运转单元。把创作拆到最底层，每一刀、每一笔、每一行字，都嵌在这个回路里：亲手操作→意外偏离→身体反打→重新判断→改向（或坚守）。这件作品之所以能一步一步往前推进，不是因为创作者事先知道终点在哪，而是因为这个回路的每一次运转，都会在当前的局面上重新生成下一步的走向。走向不是计划出来的，是被这个回路一次次逼出来的。

裁决回路的运转，不依赖任何外部担保。它不需要艺术学院的存在，不需要画廊的认可，不需要批评家的阐释，不需要艺术史的背书。一个人独自在房间里画画，没有观众，没有市场，没有体制的任何响应——裁决回路依然可以完整运转。因为他手里那支笔在纸上留下的痕迹，会反打给他自己。他自己就是那个被触动、做判断、改方向的人。这个过程在内部自我闭合，自成一体。正是在这个意义上，裁决回路构成了艺术发生的最底层的"骨架"：它是唯一不需要外部条件就能自己转起来的那台引擎。

但它的运转有一个前提：回路的各个节点，必须由同一个人亲手踩过。亲手，意味着操作环节不能外流——必须是你的手、你的眼、你的身体在承受那一下笔触的不可逆后果。亲自被反打，意味着你不能只是"看到"一个AI生成的结果，而必须是那个结果的出处与你的身体动作之间存在一条连续的因果链——是你刚才那一下手抖，造成了画布上那一小片意想不到的晕染；你退后看到它，才被它触动。亲自改向，意味着你做出的判断——无论是放弃还是沿着意外走——所改变的，是下一项具体操作的走向，而这个操作，同样要由你来亲手完成。

裁决回路这四个字，重点不在"裁决"，而在"回路"。它不是一个一次性决定，而是一个持续运转的、不被中断的闭环。裁决不是一个选择点，而是一条链。每完成一次裁决，它就会自动推进到下一环，直到作品在某一个时刻停下来。持续运转本身，就是艺术发生的方式。

## 1.1 与四条既有进路的显性对话

以上论述勾勒了本文的基本判断。但在展开之前，必须正面回应那些已经深入讨论了"过程""身体""操作"和"技术环境"的思想传统。本文的命名——"裁决回路"——是否只是一个旧概念的新说法？它切开的那个辨别面，是否已经被此前的理论覆盖了？本节将裁决回路置于四条最接近的进路之中，逐一对质，指出它究竟多出了什么。

## （一）艺术理论中的"过程"传统

坚持"过程优先于成品"的论述在艺术史上由来已久。二十世纪中期，波洛克（Jackson Pollock）的行动绘画将创作推至一个极端：绘画不再是构图的事先规划，而是画家身体在画布上运动时，"滴洒"这个行为本身所产生的不可重复的事件。批评家罗森伯格（Harold Rosenberg）曾将画布重新定义为"一个进行行动的竞技场"，而非一个承载图像的平面。这一脉思路后来被扩展为对一切行为艺术、偶发艺术和过程艺术的理论描述：创作的意义不在于最终物，而在于发生的那个无法复制的动作序列。同样，在中国传统画论中，"经营位置"的审慎与"意外生发"的随机之间始终存在一种微妙的张力。石涛（清）的"一画"论既承认一笔落纸便确立了画面的根本法则，又将后续的每一笔都视为对初始那一笔的回应与逃逸——一种在规则与偶然之间持续摆动的动态过程。

这些论述已经清晰地看到了"过程"和"意外"的生成性力量。但它们从未将这个过程拆解为一个最小可运转的闭环，从未追问：在波洛克滴下颜料的那个瞬间，他的身体到底完成了一个什么样的完整的因果循环？他的身体动作释放出颜料，颜料的轨迹与画布的相遇产生了一个他无法完全预见的形态，他的眼睛接收到这个形态，这个接收触发了他下一滴颜料的动作方向有所偏移——是这样一个不可拆分的、从操作到感知到判断再回到操作的微小闭环，在一次又一次的接续中，推动着整幅画往前走。裁决回路的过程论的区别就在于此：前者不是对"过程重要"的再次强调，而是对过程内部那台自推进引擎的精确结构的解剖。过程论看见了那台引擎的存在，但没有拆开它的缸体。

## （二）现象学中的"身体"理论

现象学传统为理解身体在知觉构成中的基础地位提供了最系统的哲学资源。尤其是梅洛-庞蒂（Maurice Merleau-Ponty）对"身体图式"与"表达操作"的分析：我们的身体不是一个被动的接收器，而是一个主动构建知觉场的原初意向系统。画家不是用眼睛"看"世界，而是用整个身体在"居住"世界——他的笔触，是他的身体对可见者的某种回应，而这种回应在他尚未意识到之前就已经发生了。现象学因此将"具身性"确立为一切创造性感知的不可还原的前提。

裁决回路接续了这一判断，但向更深处踩了一脚。现象学揭示的是：身体必定在场。裁决回路追问的是：身体的哪个瞬间，以何种因果归属的资格，在场？梅洛-庞蒂描述的是画家的身体如何在与可见世界的缠结中"延续"了知觉。但裁决回路指出，在创作过程中，有一个极其特殊的时刻——即那一下意外的反打被感知为"我刚做的事造成的"的时刻——是将身体对世界的一般性参与，切换为一条完整的、可被辨识的创作推进链的最后一道闸门。身体的在场是必要条件，但不是充分条件。充分条件是：那些由身体动作产生的物理后果，必须被同一个人认领为"我造成的"，由此触发那一下不会发生在旁观者身上的"重新判断"。这就是裁决回路对现象学"具身性"的结构化推进：它不再把身体当作一个笼统的"在场基底"，而是把身体在回路中扮演的角色分解为三个不可互换的节点——因果行动的发起者、意外反打的承受者、重新判断的执行者。现象学说"身体是一切知觉的根源"；裁决回路说"身体必须亲自踩完这三个节点，且保持同一条因果链，回路才运转"。后者是前者的一个操作性深化，而非复述。

## （三）当代媒介理论中的"操作"与"程序"

媒介理论从麦克卢汉（Marshall McLuhan）的"媒介即讯息"开始，就已经在追问技术环境如何重塑人的感知方式与行动尺度。弗鲁塞尔（Vilém Flusser）在《技术图像的宇宙》中进一步提出了一

个极有穿透力的诊断：当图像由算法生成时，人不再是在"创造"图像，而是在"编程"图像——他不再直接操作表面，而是操作程序参数。创作的手，变成了管理的手。这一脉理论正确地看到了技术对操作环境的整体性替换。在弗鲁塞尔的描述中，"编程者"站到了一个与"创作者"完全不同的位置上：他的身体不再直接触碰图像的物理表面，他的每一个决定都需经由一套符号-逻辑系统（程序语言、参数、算法）的中介才能抵达最终图像。弗鲁塞尔诊断了这种位移——他看到了旧有操作模式的整体退场，看到了人的创作姿态从"直接制造"变成了"间接设置"。

裁决回路与媒介理论的分水岭在于：媒介理论倾向于将这种替换理解为一个"环境"尺度的问题——技术重新定义了创作的生态，创作者因此需要重新适应新的条件。裁决回路则指出，技术替换的不是"环境"，而是回路上的节点。"环境"是一个弥散的、边界模糊的笼罩物，而"节点"是一个结构中的精确组件。弗鲁塞尔的"编程"已经捕捉到了创作者不再直接操作表面这一核心事实，但"编程"作为一个整体性诊断，无法进一步回答三个更精细的问题：第一，创作者从直接操作表面转向管理参数，具体是回路上的哪一个或哪几个节点被拆掉了？第二，某一个节点被外包后，因果推力链在哪个精确位置上断裂？第三，断裂的下游后果是什么——为什么上游操作的外包，会导致下游判断环节失去运转的前提？"节点外包"正是在这个意义上提供了"编程"装不下的理论收益：它不是一个更精确的同义词，而是一个解剖层次的跃迁。"编程"做的是宏观诊断——它告诉你整个创作姿态从A变成了B。"节点外包"做的是微观解剖——它拆开那条从操作到感知到判断的闭环，告诉你第一环被卸掉后，第二环为什么就收不到信号了，第三环为什么就不再被逼出判断了。这就是为什么媒介理论最终常常走向伦理呼吁——"我们应该重新重视手工艺"——而裁决回路能给出一个严格的、不依赖伦理取向的结构性诊断：不是"应该保留"，而是"如果节点外包，回路必定脱扣"。这是从"环境性诊断"到"结构因果链诊断"的跃迁。

#### （四）当代艺术界中的"人机协作"论述

近年，一批积极实践AI的艺术创作者和策展人提出了一种影响广泛的主张：AI不是工具的替代者，而是创作的"共同创作者"或"灵感伙伴"。艺术家通过与AI反复对话、修改提示词、手动调整输出图层，可以形成一个新的、人机交织的反馈循环。有策展人将这种实践称为一种新的艺术民主化——AI把创作门槛压低了，让更多人能够参与表达。这一路论述在艺评与展览文字中反复出现，其通行表述大致是：艺术家与AI之间的反复对话——输入提示、观看输出、修改提示、再观看新输出——构成了一种跨越碳基与硅基边界的创作反馈循环，人在其中的角色从"创作者"转为"回路的共同调节者"，这不是创作的终结而是创作回路的扩展。这一主张在理论上的更早根基，是玛格丽特·博登对计算创造力的经典分析：她区分了组合型、探索型与转化型创造，并论证机器至少可以在前两种意义上参与创造过程。编辑说明：本文原稿在此处曾具名引述一篇刊物评论与一段引语，编辑核验未能找到该文与该引语的出处，故改为不具名转述这一路论述的通行表述，并只保留可核验的坐标。

这类论述最典型的论证方式，就是将这种人机往返过程描述为一种"新的回写"：人修改AI的输出，喂回去，AI再输出被修改后的新版本，人再修改——这不是和画家画下一笔后退后看、然后改向一样的循环吗？

本文对"人机协作回写"论述的回应，将在第五节以完整篇幅展开。在此只需先指出其根本的混淆：它把"循环记录"与"回路运转"当成了同一件事。在琼斯所描述的那种循环中，操作者每一次修改的对象，是一个他从未亲手从不可预判的材料中接生出来的东西。他缺少那条从身体到后果的因果归属

链——因此，当AI给出新输出时，他的判断落在一条"被陈列的候选物"上，而不是"被他自己刚才那下操作逼出来的后果"上。"挑选"与"排除"，在裁决回路中是两个本体论等级不同的动作。只有后者才能触发改向。这正是本文与拥抱派论述的分歧根部。下文将证明：不是AI不能参与过程，而是AI在默认使用模式下，被设计来替代那些必须由同一个人亲手操穿每一个因果节点的环节。这个"必须"，不是美学态度，是结构条件。

通过上述四组对话，裁决回路的位置得以划定：它不是一个总结性的概念，也不是一个"综合考虑了过程、身体、技术和协作"的合成框架。它是一个切在四条进路都覆盖了、但都未曾精细解剖的那个接合部上的结构性命名——它追问的不是"过程是否重要"或"身体是否在场"，而是：从亲手操作到意外触动到重新判断的那一圈最小因果链，究竟由哪些节点构成？每一个节点被替代后，链是松了，还是断了？这个追问，是此前所有进路都未曾完成的。本文接下来的全部工作，就是让这个追问落到具体的节点上，逐一打开。

我们需要离开抽象的讨论，先落地到一个具体的现场。

## 二、一个现场：什么是裁决回路上的脱扣

2023年秋天，青年画家小杨（化名）[1] 开始画一幅以童年回忆为题材的大型油画。创作进入第三周时，她遇到了一个极度焦虑的瓶颈：画面中有一片大面积的前景，她反复尝试了不下十种方案，每一种都只画了几笔又铲掉，再画再铲。这个过程持续了整整四天。她告诉我，在那四天里，她几乎每天待在画室里超过八小时，但真正"有效"的时间加起来可能不到两小时。大部分时间，她只是坐在那里看着画布，或者站起来在房间里走来走去。

第四天晚上，她几乎就要放弃了。就在那个时候，她忽然觉得那片空荡荡的前景并不需要被"填满"——它需要的恰恰是一种被抽空的感觉。她想起了自己童年时常常一个人蹲在院子门口的空地上发呆。于是她保留了前景的一大片灰色空白，只在最边缘处用极淡的笔触画了一些几乎看不见的尘土。这个决定，事后被她自己认为是整幅画最关键的转折点。

现在，让我们把这个过程拆开，看看在那四天里，裁决回路究竟是怎么转的。

第一天，她画下第一笔前景方案——假设是一片暗色的地面。画完之后，她退后看。那一下"退后看"，就是回路的反打节点：她画下的那片暗色，实际出现在画布上，撞击了她的视觉。撞击的结果是"不对"。这个"不对"，不是一种理性分析——不是"暗色地面在构图上会抢占视觉重心"——而是一种身体性的不适感。她说："我站在那里看了几分钟，越看越难受。"这种难受，就是裁决回路上的触动。它逼着她做了一次判断：铲掉，重来。

于是她铲掉。铲掉之后，画布回复空白。但那个被铲掉的暗色并不是彻底消失了——它以"已被排除的可能"的形态，留在了她的身体记忆里。她的下一笔，已经被这个"不是"所修改了。她会下意识地回避刚才那种用色、用笔的方式，往另一个方向试探。这就是改向。

第二天，她又画下几个新方案。每一个方案都重复着同一个循环：画下→退后看→不对→铲掉→重新判断→下一个。在这个循环中，她始终没有找到一个"对"的方案。但她正在慢慢地逼近什么东西。那些一遍遍被铲掉的前景方案，并不是毫无意义的反复——每一次铲掉，都是一个"不是"的确

立。而所有那些"不是"，正在为最终那个"是"清理出一块空旷的地基。到了第四天晚上，当她的手和眼都已经极度疲惫，当她的意识不再那么紧张地想要"解决问题"，那片灰色空白几乎是自动地浮了上来。它不是从一个选项列表中选出来的——它是以唯一剩下的那个"别的方式"的形态，在她的感官里自己显现出来的。

这个过程，和任何一个AI赋能的创作过程，从根本上就属于两种不同的事。不是因为AI生成得不够好，而是因为AI的介入方式——无论是让模型一步到位地生成最终图像，还是反复迭代"手动修改→喂回AI→看生成结果"——都在结构上把裁决回路中的"亲手"这一环给请了出去。AI的生成，不需要你亲手画下那一笔。它不需要你的手在画布上产生那个无可预判的微小偏离。因此，你也就不会把你刚才做过的事情变成一个"身体反打"的对象——你只是在看一个别人替你做了的、精密的统计推断。那个统计推断不会触动你，因为你对它缺少一种特殊的因果感知——你不觉得它是"你刚才那下动作造成的"。

裁决回路的本质要求就是这一点：回路上的每一个节点，都必须是你自己踩的。它运转一次，需要四个东西：你的操作、你操作产生的意外、意外对你的反打、你对意外的重新判断。这不是一种态度，不是一种价值观，甚至不是一种选择。它是一种严苛的结构条件——一条从"做"到"感"再到"重新判断"最后回到"做"的因果链，必须不断裂，必须每一个环节都由同一颗心、同一双眼睛、同一只手亲自穿过。只要其中任何一个节点被外包，回路就脱扣。一旦脱扣，回路就不转了。不转之后，你还能产出作品——你可能产出得更多、更快、更精致——但你产出的作品，和刚才小杨在第四天晚上产出的那一笔之间，差着一条不可逾越的本体论鸿沟。那条鸿沟，叫作"这件东西是我被它推着走出来的"，而不是"我挑出来的"。

小杨并不是唯一一个感受到这种差异的人。在工作室与课堂里，类似的自述近年反复出现，其形态相当一致：说话的人并不抱怨机器做得不够好，也不主张人应该坚持手工艺；他们说的是一件更底层的事——那个曾经自动把他推进工作室的动力，在替代方案唾手可得之后，忽然失去了不言自明的效力。**编辑说明：本文原稿在此处曾引述一位知名画家的访谈原话与一份机构调查的具体数字，编辑核验未能找到该访谈与该机构，故一并删去。**这类自述目前只以零散见闻的形式存在，没有可引用的公开调查支撑它。因此在本文中，它的身份不是证据，而是一个**待检验的现象描述**：它提示的变化不发生在作品层面，而发生在"做"的欲望层面——这正是本文第七节要用一个可实施的追踪设计去证实或推翻的东西。

从这些迹象往回看，此前关于AI艺术的讨论之所以抓不住要害，正是因为它们从来不曾把问题压到回路脱扣这个层面。防御派说AI没有情感，所以他们错把回路上的"触动"当成了"情感"。拥抱派说AI可以成为灵感伙伴，所以他们错把回路上的意外反打当成了"外部建议"。策展人说AI带来了创作民主化，所以他们错把回路上的亲手操作当成了"精英技术垄断"。媒介理论家说技术重塑了创作的生态环境，所以他们以为只要讨论环境就够了，而忽略了真正被移除的不是环境，而是那个在环境内部必须由人亲自完成的最小因果链。所有这些论述，都在裁决回路的外围绕圈。而今天，AI真正切到的，是那个回路本身。

接下来的任务，就是把裁决回路的每一个节点——亲手操作、意外偏离、身体反打、重新判断、改向——逐一拆开，逐一检查它们如何在AI介入之后，从一条紧密咬合的链，变成一段可以被拆散、外包、绕过的松散模块。这个拆解，将会覆盖此前艺术理论中一直被笼统称为"氛围""过程"和"回写"的

三个层面——但不是在旧意义上使用它们，而是把它们重新锚定在裁决回路的不同节点上。氛围，是回路得以不被打断地持续运转所需要的那层免受外部监控的保护层。波谷，是回路以最慢、最不确定的速度拼命尝试运转的那个临界状态。回写，则是回路中被"意外反打"触发的那个改向瞬间的动力学本身。

下面三节，逐一展开。

### 三、氛围的另一种角色：作为回路的保护层

---

"氛围"——艺术发生的那层弥漫的、不言自明的默会信念场——在本文的框架中需要被重新定位。此前已有独立的分析[2] 指出了这一层的重要性，论证了当生成式AI将那层环绕创作活动的默会信念拆解为可识别的工序时，创作者被剥夺了那种"我正在做艺术"的理所当然的安心。这个判断在方向上是准确的，但它需要一个更精确的结构角色。氛围不是裁决回路本身。氛围不操作，不产生意外，不被反打，也不做判断。氛围的角色，是回路得以运转的"保护层"——它负责为裁决回路的持续运转屏蔽那些会打断它的外部噪声。

什么是会打断回路运转的噪声？最核心的，就是那个问题："你这样做，还算艺术吗？"这个问题本身并不产生任何作品，但它的出现会直接击穿裁决回路。因为它迫使创作者在回路运转的中途停顿下来，从"被自己的操作反打、然后重新判断"的沉浸状态中抽离，切换到一个外部审视的位置：他开始站在一个假想的观众或批评者的角度来审查自己刚才做的那一下，而不是继续让那一下的反打推动自己走向下一笔。当这个停顿发生，回路不是转不动了——它直接被挤出了运转状态。创作者的不再是继续踩回路，而是为自己辩护。而辩护，从来不会产生作品。

氛围的功能，就是在这个问题尚未被提出之前就预先消解它。它用一种弥漫在空气中的形式，不断向创作者传递一个信号："不用担心，你正在做的事情属于那个叫作艺术的范畴。"这个信号不是一份书面担保，不是一段严谨的论证，而是一种比所有这些都更底层的、渗透在每一个日常细节里的默会的确证。学院教室里的点评、展览开幕上的交谈、前辈艺术家在访谈中说起的那些曾经同样迷茫的早年经历——这些散落的碎片，共同在创作者周围形成了一层半透膜。这层膜不是用来挡住批评的，而是用来挡住那个"你还是在做艺术吗"的根本性质疑的。

这层膜之所以能发挥作用，恰恰因为它是不透明的。它允许创作者暂时不去想那些日后可能会被问到的问题——"它有市场吗？它符合当下的批评话语吗？它在历史上会被归入哪个流派？"——而只是沉浸在手头那一小块正在做的事上。这种"暂时不用想"的悬置能力，就是氛围赠予创作者的礼物。

生成式AI的暴力，便在这里使用了一种此前从未有过的方式介入。它不是批评这层氛围，不是挑战它，不是要求它扩展边界——这些历史上的先锋派运动都做过，而每一种进攻，最终都会被氛围吸收、消化，甚至反过来强化它。AI做的是另一件事：它不进攻氛围，而是绕过它。它把氛围里那些曾经只能被浸润、被默会的法则——什么画法看起来"高级"、什么配色在当下艺术界更受青睐、什么构图在历史上被尊为经典——全数拆到可见层，变成可以一键调用的参数、可以随时检视的统计规律。当这层半透膜被完全透明化，它就失去了保护功能。因为保护的前提，是你不知道外面在怎么看你。一旦你知道，你就无法不把自己放在那个被看的位置上。而一旦你把自己放在被看的位置上，裁决回路就被打断了。

这不是比喻。这是一个极其精确的心理过程。从前，一个创作者在面对空白画布时，拥有的那种"虽然不知道对不对、但感觉大概能行"的安心，其底层机制可以用裁决回路的术语重新描述：当他画下第一笔时，他的注意力完全聚焦在回路的下一个节点上——"这一笔下去，画面反馈给我什么？"他之所以能这么做，是因为那层保护膜把"别人会怎么评价这一笔"这个外部问题挡在了回路之外。现在，AI把那层膜拆了。它让他可以在任何时刻知道自己这笔下去，在"当前艺术圈品味的统计均值"上大概能得几分。于是他的注意力被从回路的内部——从"这一笔下去，画面反馈给我什么"——拉到了外部："这一笔下去，别人会觉得好吗？"这个切换一旦发生，裁决回路就停转了。他不再是被自己的操作反打的人，他是一个在替假想观众预审自己的人。而预审和裁决回路是互斥的：你不能同时沉浸在反馈闭环里，又站在闭环外面给它打分。

氛围在裁决回路中的角色因此被精确化了：它不是回路的组成部分，而是回路的运行条件。它不参与创作，但它为创作提供了一个最不可或缺的东西——一个不被监控的、可以放心沉浸在回路自身反馈中的连续时间段。它是一块无形的隔音板，隔开的是社会评价系统对创作过程每一个节点的、无孔不入的凝视。艺术史上一再被重提的那个"孤独的天才"形象，其真实的心理内核不是不需要他人认可，而是需要一个不被他人凝视所打断的完整回路运转时间。在那个时间里，他只需要面对自己亲手制造并反打回来的意外。氛围一旦被穿透，这个不被凝视的时间就消失了。回路被曝露在社会凝视下，每一个节点都可能被评论、被预测、被优化。而当回路不再是自己转动、而是被外部目光驱动的，它的输出就不再是裁决的沉淀，而是对期待的回应。这两者之间的鸿沟，就是"亲自踩出来的"与"被认可出来的"这两类作品之间那条看不见却无法逾越的裂缝。

小杨在研二上学期回顾她那段四天煎熬时，说过一句极有诊断价值的话："后来我好像再也没有像那次一样，在一幅画前面硬熬那么久了。因为我现在知道了——不是真的知道，是手里有了东西。"她说的"东西"，就是AI。她知道自己随时可以调取AI来生成临时方案。仅仅是"知道自己可以"这个事实，就已经让那层保护膜失效。因为在那个面对空白画布、不知所措的时刻，她的思绪不再停留在"这一笔怎么下去"，而是被一个元问题所介入："我为什么不直接用AI试试？"这个元问题不像直接批评那样粗暴，但它更隐蔽地切开了氛围——它让创作者在自己独处的工作室里，就开始扮演自己的评审者。氛围的穿透，在这个意义上，不需要发生在外部，它完全可以在一个创作者的头脑中自动完成。而AI的在场，恰恰提供了那个自动完成的开关。

这也就解释了为什么本文并不判定每个创作者都必须重度使用AI才受其影响。整个生态的时钟被拨快了。创作周期在缩短，展览排期在加密，市场反馈在加速。旧氛围赖以维系的那个"慢慢来也OK"的弹性，正在被一层新的大气替换——那层大气不来自任何一个画家的个人选择，而来自AI在整个艺术生态中所制造的信号密度和透明度。在那层大气里，创作者即使决定不使用AI，他也失去了曾经那种"没有人知道我在干什么、我可以在暗处摸索很久"的自由。因为他知道，他的同行们正在以他十倍的速度产出。因为画廊知道。因为社交媒体知道。因为那个曾经为他扛住这些外部压力的氛围，已经被整个环境削弱了。

裁决回路，至此失去了它的保护层。但氛围的剥落不仅仅让回路失去了遮蔽——它还带来一个更深的后果：它动摇了创作者对自己"正在做艺术"这件事的基本确信。当"这个做法在当下的品味分布中大概处于什么位置"变得随时可查，创作就不再是一个在暗处摸索的、无需自我辩护的过程，而变成了一个持续被潜在的外部目光穿透的过程。在那种穿透中，被侵蚀的不仅是回路的连续运转，还有一种更底层的身份安全感——那种"我正在走一条可以被叫作艺术的路"的不言自明的信念。这种信念的消

蚀，并不直接导致作品产出的中止，但它把创作从一种沉浸式实践变成了一个需要反复自我确认的行为。而任何需要反复自我确认的行为，都比不需要自我确认的行为消耗多得多的意志资源。这层消耗，将在下面的两节中，与回路内部的高压排除和因果归属断裂叠加在一起，共同构成发生虚无的完整动力学。

## 四、波谷的另一种本质：作为回路的极限压强

氛围为回路屏蔽外部打断，保护它不被社会凝视所穿透。但当回路内部运转遭遇极限阻力、所有反馈均为负向时，裁决回路便进入了一种全然不同的状态：波谷。[3]

波谷，指的是创作过程中那些漫长的、几乎毫无产出的枯竭期。一连几天甚至几周，画布上没有进展，稿纸上写不出一个像样的句子，排练厅里反复尝试却始终找不到动作的质感。在本文的框架里，波谷不是创作者的"情绪问题"，不是他意志力薄弱或者灵感枯竭。波谷是裁决回路以极限压强运转的那个状态。

极限压强的含义是：操作持续进行，意外持续产生，但每一次意外的反打给出的信号都是"不对"——"这一笔也不对""这个方向也不对""再来一次还是不对"。在正常情况下，裁决回路运转几次就能捕捉到一个正向反馈：你画下一笔，意外效果反而比预期更有趣，你被这个意外触动，做出改向判断，然后往新的方向推进。但在波谷期，你画下许多笔，得到的反馈全部是"不对"。回路没有断裂——你还在踩，你还在亲手操作，你还在看结果——但它陷入了一个低速、高阻、毫无进展的泥沼状态。

这种状态极度消耗耐力。因为它不再给你任何正面激励。你不再是"越做越有意思"，而是在"越做越难受"中硬撑。许多半途而废的作品就是在波谷中被放弃的。但正是在这种高阻状态下，裁决回路中所蕴含的那种"排除"机制，才能以最高效率运转。每排除一个方案——"这个不对"——回路就在物理层面和判断层面同时消去了一个可能的走法。被排除的走法越多，剩下的走法就越少。当剩下的走法少到只剩下那一个唯一可以走的方向时，裁决回路会突然加速——那个方向几乎是自动地从泥沼中浮上来，它以唯一幸存者的身份，携带被大量"不是"所筑成的必然性，直冲而出。

这就是小杨在第四天晚上经历的事。那四天里，她反复铲掉的每一笔，都是回路的一次负向运转。这些运转毫无正面产出，却在安静地做一件事：大幅缩小她以后还能走的路径空间。她在"不是"中一步步给那个最终唯一的"是"让出了路。等到第四天晚上，当她所有的"不是"几乎都已经试尽，那片灰色空白，就成了最后剩下的那条路。它不是被"想"出来的，它是被"逼"出来的。

这就是波谷的功能：它是让裁决回路的排除机制集中施压的唯一手段。没有足够的排除，你就只能在外围晃荡，永远到不了那个必然的核。排除需要时间，需要反复，需要一个人真的把每一个诱人但实质上错的方案亲自踩一遍并亲手铲掉。这个过程中，没有任何AI能替你走，因为AI的"替"在结构上与你需要的东西刚好背道而驰——你需要的是亲自走过、亲自确认"这个不行"；而AI提供的是替你生成一堆基于已有审美规律的"有可能行"的选项，让你从中挑选。挑选和排除是完全不同的两件事。挑选，发生在"你知道你要什么"的时候。而波谷期，你最无法回答的问题恰恰就是"你要什么"。你只能通过不断排除"不要什么"，来慢慢逼近那个模糊的目标。而这种排除所需要的唯一燃料，就是你亲自踩过并亲自确认无效的每一下操作。AI的生成不能替代这种操作，不是因为AI生成得不够像，而是因

为被生成出来的废稿不是"被你亲手铲掉的"——它不曾经过你的手、不曾被你退后看过、不曾让你产生那种身体性的"难受"。它对你无动于衷，你也对它无动于衷。它没有资格成为你排除列表上的一条有效记录，因为它不是你踩过的。

这就是我们理解"亲自性"这一概念最关键的鉴别标准。亲自性不是一个浪漫派的口号，不是一种对技术的道德抵抗。它是一条严苛的因果条件：只有当一项操作的全部后果——包括它的外观、它在画面上的位置、它带来的那种身体不适感——都通过你自己的身体和眼睛完整地传导过一次，并引发了你自己的判断改变时，这项操作才算被"亲自踩过"。而只有被亲自踩过的操作，才有资格成为排除链上的一环。因此，排除的有效性与排除过程是否被亲自经历之间有本质联系。亲自不是效率问题，是资格问题。

当AI填平波谷——当它在你最迷茫时给你三十个方案，让你不再需要亲自画、亲自铲、亲自一脸茫然地对着空画布——它实际上做的是釜底抽薪：它移除了排除链的燃料来源。你依然可以做判断、做挑选，但你做判断的对象不再是你自己亲手接生到这个世界上的半成品，而是一堆来自外部统计模型的候选物。你对这些候选物没有排除的资格，因为你没有亲自穿过它们。你没有资格说"这个不行"——你只能说你不喜欢它。而"不喜欢"与"它不行"之间的区别，就是裁决之浅与裁决之深之间的全部距离。

小杨后来跟我说的一句话，把这个关节钉死了。她说，那次四天的煎熬之后，每当遇到类似的瓶颈，"我的第一反应就不再是坐回去继续想，而是先看看AI能做到什么程度"。请注意这里的时态变化。四天煎熬时，她面对波谷的方式是：继续踩回路，继续排除，直到那个唯一答案被逼出来。在那之后，她面对波谷的方式变成了：先让AI给她一个答案池，她在池子里挑。这不是"慢与快"的差别，也不是"深与浅"的差别。这是裁决回路被从"高压排除模式"切换到了"低压筛选模式"。在后一种模式下，排除不会发生，因为没有任何操作是你亲自踩过的——那些操作全是AI做的，你没有资格去判断它们"不行"。你只能用"喜欢"或"不喜欢"去替换掉"对"或"不对"。这种替换的累积后果是：作品失去了必然性。它可以是好看的、精巧的、符合品味的，但它身上没有那种"它只能这样，因为它不能是别的任何样子"的重量。那种重量，是排除的负向累积在最终突破的那一刻所爆发出来的能量。你不穿越波谷，那个能量就不会来。

但这并不意味着，使用AI的艺术家就不能在另一种模式下把与AI的反复对冲过程本身做成一种新的波谷。这是一个重要的论点，它关系到本文判断的可修正边界。第七节会详细处理这一点。在此只需伏笔：如果有一个艺术家把与AI的反复碰撞本身做成一段高压的、不确定的、充满挫败感的排除过程——他不只是挑选AI生成的结果，而是一遍遍亲自覆盖它、修改它、毁掉它、再喂回去让它生成新的东西、然后再毁掉——那么在逻辑上，他就可能把AI建成了一个延长了的波谷，而不是一个波谷的替代品。这绝非不可能，但它需要满足一个苛刻的条件，后文再议。

回到当下的论证：在默认的使用模式下——也就是创作者遇到瓶颈时用AI生成一堆方案、从中挑选那个看起来最接近他想要的效果——波谷是被填平了。而填平波谷，就是切断了裁决回路中唯一能把排除效应推到极限的那个物理阶段。当回路失去了它的高压区，它依然可以运转，但它运转出的结果，将失去只有高压排除才能逼出的那笔不可替代的必然性。这就是波谷在裁决回路中扮演的角色：它是那条链上必不可少的阻力点。撤掉阻力，链还能转，但它输出的，不再是同质的裁决沉淀。

## 五、回写的另一种实质：作为回路上的意外触发的改向

裁决回路中最容易被浪漫化、也最容易被误解的一个节点，就是"回写"。关于创作中的回写机制，此前已有独立的分析：画家画下一笔，退后两步看，刚才那一笔改变了他对整幅画的感受，这种改变反过来指引他画下一笔。这个判断是准确的，但本文需要给它一个更精密的承重面。回写不是"操作之后的那一下感悟"，而是裁决回路中特定两个节点之间的触发器：它发生在"意外偏离"和"重新判断"之间。当创作中的一次亲手操作产生出操作者本人没有预料到的效果时，这个偏差并不直接被转入下一项操作，而是先经过一道闸门——它必须被感知为"触动了"操作者，才可能激活重新判断。而被感知为"触动"，不是一种美学评判，而是操作者的身体和感知系统在瞬间自行进行的反应。它与操作者在那一瞬间是否有理论框架、是否有职业素养、是否有艺术家身份都无关。它与一件事有关：这个东西是刚刚被他自己亲手制造出来的。

### 5.1 因果归属：触动的前提条件

这也就是为什么回写在AI介入后出现断裂，不是发生在"AI不能产生意外"这一层。AI当然能产生意外。任何用过生成模型的人都知道，输入同样的提示词，隔十秒钟再生成一次，结果可能截然不同。AI的输出充满了不可预见性。但这种不可预见性，与手抖、颜料多蘸了一点、力道重了一点点所造成的意外，属于两个本体论层级。差别不在"意外的外观"上，而在"意外与操作者之间的因果归属"上。一个创作者面对画布上的那片意外的晕染时，他知道："这是我刚才那一下手抖造成的。"这个"我造成的"是不可撤销的因果事实。他被这片晕染触动，不仅仅因为它看起来有趣，而更因为他与它之间存在着一条不隔的、连续的所有权链——这片东西是他身体行为的直接物理后果。他不只是看到了它，他是认领了它。

这个认领过程，不是一个理性归因的、事后的推理。它是即时的、预反思的。被这片意外晕染触动，本质上是一种极短时间内的因果感知，而不是一种审美判断。要理解这一点，可以借用人对自身动作之结果的"能动性"（sense of agency）这一认知结构——当一个行动者执行一项动作并立即感知到其后果时，他会自动地、不需要任何概念中介地将那个后果注册为"我的动作造成的"。这种注册会引发一系列即时的生理与注意层面的反应：心率微调、瞳孔变化、注意力的瞬间集中。这并不是说创作者在"分析"这片晕染是否有趣，而是他的身体在几百分之一秒内已经做出了"这是我自己制造出来的意外"的标记。这个标记，就是触动的生理基础。触动不是一种高级情感，而是一种感觉运动层面的因果再认。[4]

而AI生成的意外，无论多么精彩，不携带这种因果再认的前提。它不是操作者身体行为的因果后果。它是统计模型在海量数据中重组出的一个概率峰值。操作者可以看到它，可以欣赏它，可以选择它，但他无法在感觉运动层面上将它注册为"我的动作造成的"。因为在他和这个意外之间，不存在一条连续的、从肉身经由物理材料到最终效果的因果链。即使他反复修改提示词、调整参数、甚至手动在AI输出的基础上用数字画笔做出局部干预——他和最终那个结果之间的因果关联，仍然是离散的、跳跃的、被黑箱隔断的。他知道他"做了什么事情导致了某种输出"，但他不能像感知自己的手抖造成的晕染那样，在身体层面追踪到那个结果是如何从刚才那一个具体动作中生长出来的。布与手指之间那不到一毫米的不可预判的偏差，在这里被代之以提示词和扩散模型之间那数百万维度的潜在空间运算——后者在因果上是完全不透明的。缺失了因果链条的连续追踪，AI输出的那个意外对他而言，是

一种被呈交上来的陈列品，而不是一个刚刚从他的身体里长出来的怪胎。它触发了"好感"或"惊喜"，但无法触发改向所必需的那种"我刚做的事把我自己往一个没预料到的方向推了一把"的身体性认知。

这个差别极其微妙却极其要紧。回写的实质，不是意外被看见，而是意外被认领。而认领的前提，就是回路上的"亲手操作"节点与"意外偏离"节点之间没有断裂，且操作者能够以连续感知追踪从动作到后果的整条因果链。一旦这一环外包，意外偏离就不再是回路内部产生的，而变成了在回路外、以第三方的形态递交给操作者的。它可以被浏览，却不能被认领。不能被认领，就不能触发重新判断。这就是AI抽掉了回写的节点，不是因为它拿走了意外，而是因为它切断了意外与操作者之间的因果归属，从而让意外变成了一种陈列品，而不是回路上一环。

从能动感的角度看，AI所替代的不仅仅是操作的行为本身，而是能动感的深层源头——那种"是我造成这个意外"的因果归属体验。认知科学和现象学的研究区分了不同层次的能动感：低层次的是对单个动作后果的即时感觉（"我推了一下，它动了"），高层次的是对一系列动作的整体归属感（"这幅画是我画的"）。裁决回路所要求的，恰恰是低层次能动感在每一个操作节点上的持续累积——每一次"这一笔是我画的，这一片晕染是我造成的"，都在为后续的重新判断提供不可替代的因果燃料。AI替代的不只是操作的执行，而是这整个因果归属链在身体的感知运动层面的持续生成。这意味着，"回路脱扣→创作动力缺失"这一链条，不只是一个哲学推论，而是有认知科学上的具体机制作为支撑：当反馈闭环（意外反打→触动→判断）被打破，创作者失去了内在奖励机制——那种由意外反打带来的认知刺激与惊喜感——导致心理学家称之为"心流"（flow）的沉浸状态难以维持。创作动力不是一种抽象的意志品质，而是一个需要被每一圈回路的成功运转反复喂养的生理-认知过程。回路一脱扣，那个喂养就停了。动力不是在中断之后"减弱"了，而是在中断的那一刻就失去了它的再生条件。

## 5.2 "人机协作回写"为何不构成回路

这层结构也使我们得以看清一种广为流传的"人机协作回写"论述为什么站不住脚。该论述大致是这样：人可以反复手动修改AI生成的图像，再喂回去让AI继续生成，手动修改与AI生成交替多次，形成一个反馈循环——"这不就是新的回写吗？"

这不构成回写，因为那个循环的每一个环节里，操作者的判断落在不同的对象上。当他对着AI生成的图像做手动修改时，他是在修改一个他不是作者的东西。他修改完了，喂回去，AI生成了新图像，那个新图像依然不是他亲手从那片不可预判的材料中逼出来的。他与那个图像之间，没有那条"是我刚才那下操作导致了它"的因果链。他有的，是一长串循环记录。但这种循环记录不论反复多少次，都无法等价为裁决回路中那一下不可撤销的、属于自己的意外触动。因为它缺少因果归属。这一个词，就是回写的全部重心。

具体而言，在前述琼斯所描述的那种"与机器的回写"中，艺术家的工作流程是这样的：他坐在屏幕前，输入一段提示词，AI输出四张图像。他浏览这四张图像，选择其中一张，在Photoshop中做出局部调整——改变了某个区域的色调，涂抹了一处他不满意的细节。然后他把修改后的图像喂回给AI，要求它"在此基础上再生成"。AI再次输出结果。他再次修改。如此反复多次。

这个流程在表面上看，似乎构成了某种"循环"：他做出操作，AI响应，他再操作，AI再响应。但这里缺少裁决回路中最关键的那个东西：他的每一次手动修改，不是被"刚才那下操作所产生的、不可预判的意外"所触发的。他没有"亲笔画下一笔、退后看、被那一笔的意外效果反打"的经验。他只是在

管理一台替他生成图像的机器，而他对这台机器输出的结果，没有因果归属——他不能指着AI生成的某一片纹理说："这是我刚才那下鼠标拖曳的直接物理后果。"他能做的是挑选、修改、再喂回。这三个动作——挑选、修改、再喂回——在因果归属的层面仍然是离散的，而不是连续的。挑选的对象不是"我刚才做的"，修改的基底不是"我亲手从材料中接生出来的"，再喂回的后果不是"那下修改的直接回弹"。

这就是为什么"循环记录"与"回路运转"是两种完全不同的事。循环记录是一条链条上反复留下痕迹，但每一个痕迹与其前一个痕迹之间，没有因果归属的连续性。回路运转则要求每一个节点的输入，都是上一个节点的输出的因果后果——不是逻辑上的后果，不是程序上的后果，而是物理上的、你自己的身体动作在不可预判的材料上产生的那一下不可撤销的偏移所造成的后果。这二者之间的差别，不是一个"程度"问题，而是一个"种类"问题。没有人能通过增加循环记录的数量，来跨过这道种类的鸿沟。

### 5.3 延迟的功能：暴露因果链的时间间隙

这同样解释了"延迟"在回写机制中为何不可消除。许多讨论已经指出了身体在创作中的角色——手会累，眼睛需要适应画布的尺寸，油彩不会立刻干，凿下的石头无法长回去。这些物理的、生理的限制，被本文统一收归到裁决回路的同一个功能上：它们为"亲手操作→产生意外→反打→被触动"这个序列提供了不可压缩的时间间隙。在这个间隙里，操作者的感知系统不是静止的，它正在消化上一笔的后果，并且正在为下一笔预备姿态。更重要的是，这个间隙将因果链暴露在操作者的身体感知中：在颜料未干的那几个小时里，那个意外的后果就是他刚才那一笔仍在起效的物理证据。他可以一直看着它，一直消化它，一直让它在他身体里持续产生微弱的效应。如果这一步被压缩成秒级的调整参数、点击生成，那间隙就会消失。间隙消失，因果链就被时间性地上了最短路径化——操作者失去了感知因果关联所需要的物理时间的跨度。AI作为精度机器，其设计初衷本身就包含了消除延迟、撤销不可逆性、用多版本并行来替代单次冒险。这三点分别对应于裁决回路的三个核心特征——不可压缩的时间间隙、不可撤回的物理后果、被限制在单通道上的人体认知——被系统性地优化掉了。这不是AI的缺陷，而是AI的使命。它被设计出来，就是要用来消除这些"低效"的。

因此，把AI与传统的画笔或照相机相比，在裁决回路的视角下失效了。照相机依然需要人亲自选择角度、等待光线、在暗房里亲手操作化学药剂——这些环节中的每一个，都仍然被一条完整的裁决回路所穿过。它不是代偿回路，它是延长了回路——把决策点从画布前摊到了暗房里。而AI在默认模式下走的是相反的路：它把回路上的环节一个个卸掉，把它们变成独立的、可以由模型自动完成的功能模块。这二者之间的差别不是程度，而是方向。照相机让人以另一种方式继续踩回路；AI则被设计来让人不必再踩。而一旦你不踩，你就不能再说那个意外是你自己撞出来的。这就不是一个人在和工具较劲，而是一个人选择了让渡回路上一环。

## 六、回路断裂之后：发生虚无——一种结构后果及其生成

当裁决回路的保护层被穿透（氛围透明化），当它的高压排除区被填平（波谷代偿），当它最核心的改向节点被切断因果归属（回写外包），回路本身便在创作者对作品的日常推进过程中，一段接一段地脱扣。这不是"创作变得更难了"，而是创作这件事在日常发生的意义上，失去了它内在的推动

力。它不能再靠回路每一圈的自行运转来产生下一步的走向。创作者不能再像从前那样，只是坐在画布前，凭着一项接一项的亲自操作，被材料反打、被意外触动、被排除逼着改向，就这么一步一步走下去。他现在面对的是：回路不转了。他必须为每一步寻找一个外在的理由。

这就是本文所称的"发生虚无"。必须立刻澄清：它不是一种心理状态、情绪体验或意志力衰退。它是一个结构性的动力学后果，更精确地说，是一个因果动力来源被切断后，创作推进陷入空转的状态。裁决回路在正常运转时，它的每一个节点都向下一个节点输出一道因果推力：操作产生意外，意外通过因果归属触发反打，反打逼出判断，判断指令下一项操作。这条因果推力链是自足的——它不需要任何外源性的"我想继续""我应该继续""继续是有意义的"这类动机来充当燃料。回路的每一次自转，自动将"为什么要画下一笔"这个问题的答案，嵌在了上一笔的后果对创作者发出的回应中："因为刚才那一下还没消化完。"在这个意义上，创作者的持续推进不是自律的结果，不是激情的产物，更不是对艺术使命的觉悟。它纯粹是回路绕着链一圈一圈往前推的自动后果。

AI介入之后，回路节点被外包。外包一个节点，那条因果推力链就在那个位置断开。上一笔的后果不再反打给操作者——因为它不是操作者做的。意外不再被认领——因为它没有因果归属。判断不再被逼出来——因为没有东西来触动它。于是，那条曾经自动输出"继续"指令的链条，停了。这就是发生虚无的结构机制：它不是"我觉得空"的心理感受，而是创作推进的因果动力来源被切断了。后续操作不再由回路自生成，而必须依赖外源性的理由注入——创作者必须主动去找一个"我为什么要继续"的答案，来填补那条因果链断开后留下的推力真空。这个状态可以完全不带任何情绪色彩：创作者可以心态积极、精力充沛、对艺术充满信念——但如果回路的因果链断了，他仍然需要一次又一次地、在没有任何内在反打催促的前提下，自己为自己下达推进指令。这种"自己给自己下指令"的动作，和"被回路推着走"的动作，在结构上是两种完全不同的事。

裁决回路未断裂时，"继续"不是被决定的，而是被执行的——回路运转本身就不给你停下来的空档。回路断裂之后，"继续"变成了一个必须被反复做出的裁决。每一次裁决，都是一个独立的、需要消耗认知与意志资源的行为。而且，它不再是回路上的裁决——不是那种被上一笔的后果逼着做出的、有明确对象("这个不对，铲掉")的判断——而是一种元裁决：裁决自己是否还要站在画布前。这种元裁决没有回路。它做完了就做完了，它无法自动触发下一环。你裁决了"继续"，但你仍然需要在此后的数小时里，在没有被任何意外反打的状态下，硬生生地画下第一笔、第二笔。而画下的每一笔，如果没有经过亲手操作与因果归属的完整闭合，就不会生成新的反打，不会推动下一笔。你裁决了继续，但你并没有重新启动回路——你只是给自己发了一张许可，而许可不是燃料。这就是为什么发生虚无不以"做不出作品"为首要表现，而以"不想做"为表现。不想做，不是懒，是在因果动力缺失之后，每一次"做"都需要一次额外的意志注入，而这种注入的能耗是指数级的。

但到这里，论证还不能停。只把"不想做"描述为结构后果，仍然是一种发现学的姿势——它把"发生虚无"当作了一个因变量，一个被回路脱扣"导致"的现象。发生学的追问必须再踩一步：当这个状态成为常态，当创作者日复一日地面对"必须手动为自己注入理由"的处境时，这种处境本身，是否正在生成某种新的东西？旧结构断裂后，新结构是否正在从裂缝里长出来？

这个问题，本文不打算给出完整的回答，因为此刻它更多的是一道需要被持续观察的发生学切口，而不是一个可以被盖棺定论的判断。但至少可以描述一个正在浮现的方向：当回路不再自动运转，创作者被迫将"是否继续"作为一个需要反复裁决的元问题来面对时，他与创作之间的关系，正在

发生一种不易察觉但性质重大的位移。他不再是被回路"推着走"的人，而是变成了一个在回路的每一个节点前面停下来、问自己"这一环，我要亲自踩吗"的人。持续推进不再是一件不需要理由的事，而变成了一件需要理由管理的事。

这种位移，正在将创作者从回路的"内部操作者"切换为回路的"外部管理者"。从前，他站在回路里，他的注意力完全被正在执行的那一下操作和它即将产生的反打所占据。回路本身是他行动的媒介，而不是他审视的对象。现在，回路脱扣了，它不再是自己转的，而是一段需要被手动拼接的松散环节。于是创作者不得不退后一步，从"站在回路里"切换到"站在回路外"，看着那条曾经连贯的链，决定哪一段自己踩、哪一段外包、哪一段放弃。他不再是一个沉浸在操作-反打-判断循环中的行动者，而是一个管理着一条不再自转的管线的人——一条由他自己、AI、材料、时间、市场预期拼接而成的创作管线。

这种"管理者主体性"——姑且这样暂命名——并不是一个贬义词。它不等同于"堕落"或"浅薄"。它只是在描述一个正在生成中的事实：裁决回路脱扣之后，创作主体性本身正在经历一次重组。旧的主体性是被回路推着走的——它不需要理由，它只需要继续。新的主体性则必须学会在理由空转的状态下维持创作，必须学会管理自己的意志资源、分配自己的亲自性、在"踩"与"不踩"之间做出一个个孤立的裁决。这种主体性不是"被剥夺了回路的人"，而是"正在学习如何在回路不转的情况下继续产出作品的人"。它会带来什么？会产出什么样的作品？那些作品身上会带有怎样一种不同于"必然性重量"的质地？这些问题，正是裁决回路脱扣之后，艺术发生学下一步需要追问的问题。

小杨在第四次煎熬之后的那个细微变化，正是这种主体性重组的最初征兆。她说她后来的第一反应不再是坐回去继续想，而是先看看AI能做些什么。这不是她变懒了。这是她创作时的内在推动器——那个以前会自动转起来的回路——在某几个关口上被绕过去了。她使用的不是重度AI介入，只是偶尔调取候选项；但候选项的存在本身，就已经在因果动力链上制造了一个可能的岔路口。那个岔路口的存在，削弱了"必须自己踩"这一结构条件所施加的强制力。她面对画布的时间变少了，不是因为她在逃避，而是因为她与画布之间的那条强迫症式的因果链变松了。

这或许是目前最被低估的后果。当前关于AI与艺术的公共讨论，多集中在产出端：AI能生成什么？人还能不能做出AI做不出的作品？但发生虚无不在产出端，它在开端之前。它发生在那个每天清晨走进工作室、面对空画布、手还没有抬起来的间隙里。那个间隙，从前是不需要填充的——回路会接手。现在，它需要一个理由来填充。当你知道，你不拿起笔，AI也可以替你完成一个看上去差不多的东西时，"为什么不拿起"就变成了一个需要被抵抗的东西。而抵抗，和回路运转不是同一回事。回路运转是顺流而下，抵抗是逆流而上。顺流不需要理由，逆流需要理由。而大多数发生，都是在不需要理由的阶段被一点点推出来的。一旦需要理由，发生就变了质地。

## 七、反驳与边界：在何处这个判断会被修正

---

以上论证层层推进，最后归结为一个核心命题：裁决回路是艺术发生最小可运转单元；AI的系统性介入，是通过为其回路上的亲手操作、意外偏离、身体反打、重新判断各节点提供替代方案，导致回路脱扣、发生虚无这一结构后果降临。这个命题不依赖特定的艺术史叙述，不依赖对"艺术"的定义之争，不依赖对"创造力"本质的假定。它只依赖一个结构条件——那条从亲手操作到重新判断、必须由同一颗心、同一双眼睛、同一只手完整穿过、中间不曾外包给任何非人代理的因果链——是否保持

完整。但它毕竟是可错的。本节主动处理三个最能动摇本判断的质疑，并设置可使本判断被修正的检验条件。

第一个质疑：如果裁决回路真的如此根本，为什么那么多使用AI的艺术家并没有坠入发生虚无？

这需要严格区分两种截然不同的状态。大量艺术家在使用AI时，并不依赖AI来替代裁决回路中的操作节点。他们用AI做前期调研、搜集视觉素材、快速生成参考图来拓宽自己的想象范围。但当他们站到画布前、拿起画笔时，那一整套裁决回路仍然是他们自己在踩——亲手操作、亲自被反打、亲自改向。AI对他们而言，是一架高级的白板前端的参考图生成器，并没有插进回路的闭环中。本文的判断从来不反对这种用法，也不是对它提出批评。

另一些艺术家则走得更远：他们把AI完全嵌入到裁决回路里，让AI替他们完成回路上的关键节点。这些人——尤其是在学生和新晋创作者群体中——恰恰是发生虚无最密集降临的人群。大量的访谈数据和工作室观察能够支撑这一判断，虽然尚缺系统性的实证研究。但现象已经清晰可辨：越是用AI替代操作节点的创作者，越容易在长期维度上出现"不想进工作室"的倾向。他们不是做不出作品，他们是不想做。不做，就是因果动力缺失的结构后果在行为层面的表现。需要如实说明：这一观察目前只有工作室层面的零散见闻可依，本文在编辑核验中没有找到可供引用的调查数据或公开访谈来支撑它。因此它在本文中的身份是一个**待检验的预判**，而不是已有证据的结论——下面给出的追踪设计，正是为了让这个预判可以被证实或推翻。

本文的判断在此一端的预判是：随着AI工具的进一步易用化和深度植入，出现这类倾向的人群会持续扩大。这个预判是可检验的：如果未来五年的大规模追踪数据显示深度使用AI的创作者的"创作持续周期"不比亲力亲为的创作者显著更短——且排除了专业水平、工作条件等混淆变量——那么本文对回路脱扣导致因果动力缺失的因果判断就需要重审。具体而言，可设计为：追踪两组水平相当、工作条件相似的创作者（一组深度使用AI代偿回路节点，一组亲力亲为），在五年内记录他们的创作频率、持续周期、自我报告的创作动力变化，以及独立评估者对其作品可持续产出能力的判断。若深度使用AI组的各项指标不低于对照组，本文的脱扣假说则被证伪。

第二个质疑：会不会有创作者把使用AI本身做成一种新的裁决回路——与AI的反复碰撞变成了一个更漫长的、充满意外与挫败的排除过程？

本文将此质疑视为最重要的可修正窗口，并在此留下证伪接口。在结构上，这是可能的。可以设想一个创作者，他不只是用AI生成方案然后挑选，而是把AI的每一次输出都当作"必须被亲手毁掉、覆盖、重新喂回AI让它生成新东西"的原料。每一次毁掉与重新喂回之间的那个"不对"，都由他的眼睛和手亲自确认，并且直接触发他对下一轮输入的重新判断。在这种模式下，AI不是回路的替代者，而是被嵌入了回路——它成了那个"意外偏离"的放大器：创作者输入的东西与他最终得到的东西之间的巨大跳跃，替代了传统中画布上那一点点微小跑偏的不可预判性。而创作者与那个跳跃出来的结果之间，有一条因果归属链：这个结果是他刚才那一下毁掉和修改逼出来的。他认领它，因为它就是那一下逼出来的。

这种模式如果能够被稳定地、大规模地实施——不是一两个特例，而是在不同媒材、不同背景、不同年龄段的创作者中都出现了可被识别为同类现象的作品群——那么本文"AI系统性地导致裁决回路脱扣"的判断就需要被大幅修正。不修改的核心判断不是回路本身不成立，而是AI必然会脱扣回路这个

具体断言。修正后的版本将是：裁决回路依然不可外包，但AI可以被嵌入回路内部，成为回路某个节点的增压装置，而不是代偿装置。

但这个修正需要满足的条件极其苛刻。一个真正被嵌入裁决回路内部的AI使用方式，必须同时满足三个条件：（1）每一次AI输出，都必须被操作者亲手进行一项不可撤回的物理干预（覆盖、铲掉、烧毁、重新绘制其上），而不是仅仅在屏幕上做图层调整或参数修改；因为只有物理干预才能在那条因果链上留下不可抹消的"我做的"印记。（2）每一次干预的后续回弹——也就是对AI下一次输出的影响——必须由操作者本人承受全部认知不协调的压强，而不能简单跳过；也就是说，AI再输出一个东西时，操作者必须亲自把那个东西与他刚才所做的物理干预放在一起看，让两者之间的张力反打回来。（3）最终作品中必须残留那些被亲手毁掉的层面的痕迹，且这些痕迹在观看中可被独立于作品自述文本而感知——不能是观念上的"知道"，而必须是视觉上的"感觉到"。这三个条件同时满足时，我们就不是在讨论一个挑选者，而是在讨论一个仍然完整站在裁决回路里、把AI建成了他自己那条高压排除链上的一个增压站的创作者。反之，如果只满足其中一二，那仍然是在以某种方式让渡回路节点，只不过方式更精巧而已。

本文的证伪条件因此可以精确表述为：如果能够大规模出现同时满足上述三条件的AI深度使用艺术家群体，且其作品在独立观看中被识别的"亲自性重量"——由完整因果归属链所产生的、可被观者感知的作品"压实感"或"必然性"的强度——不低于同等水平的手工创作者，则本文关于AI系统性地拆解裁决回路的判断需要被撤回。[5]

第三个质疑：本文是否把裁决回路过于实体化了——它难道不可能是一个想象出来的、被文化建构出来的结构？

裁决回路不是任何意义上的"自然事实"。它不是生理学的断言，不是神经科学的假说，更不是对某种普世艺术本质的宣称。它是一个操作性的结构切分：只要你承认，在某种创作活动中，存在"操作—意外触动—重新判断—改向"这样一圈可被经验识别的最小闭环——且这个闭环在历史上被无数艺术家的创作自述反复描述过——那么你就可以把它拆出来，作为一个分析单元使用。本文所做的全部工作，就是把这个闭环的结构条件厘清，然后检视AI是如何系统性地为每一个节点提供了一种代办选项。这个操作是否成立，不取决于是否存在"人类的普遍创作本能"，而只取决于论文是否给出了这个闭环足够清晰的边界的操作定义。前文已经做出：亲手操作、意外偏离、身体反打、重新判断、改向——五个节点，一个因果归属。如果后续研究能证明，这个闭环可以被进一步拆解——比如，"身体反打"本身足以被AI模拟，而创作者仍然能被它触发真正的改向——那么本判断中关于"AI不能替代意外触动"的那一部分就需要被修正。但如果那个"被它触发"的前提，依然是创作者知道自己亲手做出了那项操作，那么因果归属仍然没有被打破，回路就没有被外包。那将是本文判断在更深层面被印证，而不是被推翻。这是留给未来五年创作实践和实证研究的裁决。

## 八、结论：回路上的再踩

本文的论证始于一个观察：关于AI与艺术的讨论，绝大多数回荡在"作品"或"能力"的层面，而那道最深的裂缝，裂在更下面——裂在艺术发生过程中，那条从亲手操作到意外触动到重新判断的因果链还能不能保持连续。这个因果链，本文称之为裁决回路，并论证它是艺术发生最小可运转单元。它的运转不需要外部担保——不需要体制、市场、批评史的托举——但它的运转需要一个前提：回路上的

每一个节点，都必须被同一个人的身体和意识完整地穿过。AI不是不能"创作"，而是在默认使用模式下，被设计来为裁决回路的各个节点提供不必亲自踩过的替代方案。这种替代的累积，并不导致创作的终结，但导致回路的脱扣。脱扣之后，创作不再能在惯性的自驱动中推进，它必须在每一步之前被重新裁决。而发生虚无，就是在这个因果动力来源被切断、后续推进必须依赖外源性理由注入的状态中产生的。

本文没有建议任何"应该怎样使用AI"的规范。它只是指出了一条结构边界：如果你让自己亲手来踩裁决回路的那条链上的每一个环节——如果你不把你的手、你的眼、你的身体反应、你的重新判断外放给一个模型——那么你就还站在艺术发生的内核上，无论你是否同时在用AI做辅助。如果你把回路上的某一个节点让渡了——哪怕只是一个——那么从那一点开始，回路就开始脱扣。你依然能制作出高完成度的作品，但它不再被那条因果链推动着向前走。裁决回路的脱扣和因果动力缺失的到来，与你产出作品的品质不直接相关。但它们与一件事直接相关：你还能不能在长时段里，日复一日地，在不询问"为什么还要画下一笔"的情况下，走进工作室，拿起笔，画下去。

但故事并没有在这里结束。因为发生学不负责哀悼，它只负责描述新结构如何从旧结构的裂缝里长出来。本文在第六节末尾指出了正在浮现的方向：当回路不再自动运转，当"要不要继续"变成一个需要被反复裁决的问题，创作者被迫退后一步，从回路的内部操作者，变成了回路的外部管理者。他不是不再创作，而是开始管理一条由他自己、AI、材料、时间、市场预期拼接而成的创作管线。这种"管理者主体性"——如果它确实正在生成——不是裁决回路的复活，而是裁决回路脱扣之后，创作主体性正在经历的一次重组。这重组会带来什么样的作品？那些作品会带有怎样一种不同于"必然性重量"的质地？它们会在观看中触发怎样一种不同于"被作者的亲手性压住"的经验？这些问题，本文无法回答，但它们恰恰是裁决回路这个判断打开的新问题域。一个发生学的判断，不以句号结束，而以逗号推进。

本文在第七节留下了两个可以被经验数据撼动的窗口。如果未来出现大批完整保留裁决回路三个苛刻条件、却深度使用AI的创作者群落，且他们产出的作品在独立观看中能被识别出同等"亲自性重量"，那么本文关于AI系统性地导致回路脱扣的具体因果断言将被修正。如果未来的实证研究能够证明，被AI训练出来的新一代创作者所经历的因果动力缺失并不比手工创作者更显著——且这种对比排除了专业水平和选择效应的干扰——那么脱扣导致发生虚无的那条因果链也需要被重新检视。在这些被触发之前，本文的判断站在这里。

最后一个问题：裁决回路还必须被保留吗？我们是否可能在有朝一日，活在一个完全不同的世界里，在那个世界里，"亲自操作—意外触动—重新判断—改向"这条链的确不再被视为艺术发生所必须的东西？人们在那个世界里做出许多精美的、动人的、深刻的、甚至能改变人看待世界方式的作品，但它们身上的每一个节点都可以被代偿，而不影响其最终效果。

也许吧。但那将不再是本文所描述的那个世界。在那一天到来之前，当回路的节点被逐环让渡，当踩回路本身不再被这个生态默认为是创作发生的必要条件，艺术的发生机制本身将经历一次重组——不是回退到某个想象中的手工黄金时代，而是变异为一套新的配置。在那套配置里，"亲自"不再是不言自明的前提，而变成了一个必须在每一次进入工作室时被重新裁决的选项。创作者不再是"被回路推着走的人"，而是"在每一个节点前停下来、决定这一环要不要自己踩的人"。两种身份之间，差的不是艺术的好坏，而是艺术发生的条件本身是否已经被改写。

## 注释

---

[1] 材料说明：本文对青年艺术家小杨（化名）的追踪为作者于2023至2024年间田野调查的一部分。案例中的对话与创作细节经当事人同意用于学术分析，并已作匿名化处理。该案例旨在为本文的理论框架提供一个思想拓展性质的例示，而非经过同行评议的实证数据。为突出理论焦点，案例中的部分场景与对话经过了简化和合成处理。

[2] 关于氛围、回写等概念的先期独立分析，出自作者此前未发表的研究手稿。本文将这些概念重新锚定于裁决回路的各节点之上，赋予其更精确的结构位置，而非对旧分析的重复。

[3] "波谷"是本文提出的一个技术性概念，特指裁决回路在极限压强下持续输出负向反馈（"不对"）的创作阶段。它不同于通常所谓的"创作瓶颈"或"灵感枯竭"，后者多被归因于心理或情绪状态，而波谷所描述的是回路中排除机制的动力学高压状态，与创作者的意志品质无直接关系。

[4] 关于能动感与因果归属的认知机制，可参见Shaun Gallagher, *How the Body Shapes the Mind* (Oxford University Press, 2005) 中关于"身体图式"与行动感知的讨论；以及心理学家米哈里·契克森米哈伊 (Mihaly Csikszentmihalyi) 关于"心流"状态中即时反馈闭环的核心作用的经典论述 (Flow: The Psychology of Optimal Experience, Harper & Row, 1990)。本文在此并非严格援引某一特定实验结论，而是借其描述裁决回路中"认领"动作的预反思特征以及反馈闭环对内在动机的支撑机制。

[5] 关于可证伪条件的操作化建议：所谓"大规模出现"可理解为在多个独立创作社群、媒材领域及地理区域中均能被观察到；所谓"亲自性重量"的识别，可设计为双盲实验，由专家与非专家观者在完全不了解创作过程的情况下，对作品的"必然性"或"压实感"进行评分。若未来此类经验证据表明AI深度使用组的评分不低于同等水平的手工创作者组，则本文关于AI系统性地导致回路脱扣的因果判断需要被修正。

## 附论一·最近邻切割：裁决回路不是什么

---

一个新命名要站住，必须先经受住那些"这不就是某某说过的吗"的质问。本节请出三位离裁决回路最近对手，逐一划出可裁决的分离线。第一位是最危险的，因为它几乎逐字描述了本文的闭环。

### （一）与舍恩「情境的对话」的切割：中介得了，还是中介不了

唐纳德·舍恩在《反思性实践者》里描述过一个几乎同形的过程：实践者对材料做出一个动作，情境会"回话" (back-talk) ——它以超出预期的方式回应，实践者据此重新框定问题，再做下一个动作。设计师在草图上画一笔，草图回话；教师抛出一个问题，课堂回话。这条"动作—回话—再框定—再动作"的链条，与本文的"亲手操作—意外偏离—反打—重新判断—改向"在形态上高度重合。如果两者可以互相替换，本文就只是把舍恩的话换了一套词。

分离线在回话能否被中介。舍恩的回话是一种认识论事件：它可以经由语言、图纸、他人转述、乃至一份分析报告抵达实践者，因为它要触发的是"重新框定问题"这一认知动作。一个设计师完全可以通过助手的口头描述得知"这个方案在现场行不通"，从而完成一次合格的再框定；一个教师可以通

过课后问卷得知课堂的对话。舍恩的框架不要求实践者的身体在场，只要求他保持反思。本文的反打不是这样：它要触发的不是重新框定，而是**触动**；而触动的前提条件是因果归属在感觉运动层面上的即时注册——“这是我刚才那一下手抖造成的”。这个注册不接受任何转述。助手告诉你“那片晕染很有意思”，与你自己看见你的手造出了那片晕染，在裁决回路上不是同一件事。

由此得到一条可裁决的分离线：**设置一个“高回话、零亲手”的条件**——让创作者不动手，由一位训练有素的助手执行全部操作，并把每一次意外的效果完整、即时、忠实地描述或展示给他，由他做出全部判断与改向指令。舍恩的框架预测：只要回话通道足够高保真，反思性对话可以照常进行，作品的质量与实践者的学习不会系统性受损。本文预测：这一格里回路是断的——判断仍可做出、作品仍可完成，但创作者在长期维度上出现与深度使用 AI 者相同的动力衰减，因为他从未在感觉运动层面认领过任何一个意外。这一格是可实施的（口述创作、助理执行在当代艺术中本就是常见工作方式），因此这条分离线不是思辨的。

## （二）与英戈尔德「与材料的应答」及物质介入理论的切割：连续对话，还是闸门

蒂姆·英戈尔德的《Making》主张，制作不是把预想的形式强加于惰性材料，而是匠人与材料力量之间的一场“应答”（correspondence）——顺着木纹走，跟着黏土的湿度走。拉姆布罗斯·马拉福里斯的物质介入理论走得更远：认知不是发生在头脑里再输出到物上，而是分布在手、工具与材料构成的回路中，陶轮上的泥与陶工的手共同构成了那个思考着的系统。两者都在说“材料参与了创作”，与本文的反打节点邻接。

分离线在**有没有那道闸门**。英戈尔德与马拉福里斯描述的是一种**连续的、无门槛的耦合**——材料时时刻刻都在参与，手与泥之间不存在“这一次算数、那一次不算数”的筛选。本文的裁决回路则要求一道明确的闸门：偏离必须被感知为“触动了”，才进入重新判断；没触动的偏离会被直接吸收进下一次操作，回路空转。这道闸门解释了一件应答理论解释不了的事：为什么同一个创作者在同一批材料上，有些天回路转得飞快，有些天做了一整天却什么也没被触动。应答是连续的，触动是离散的。

可裁决差别：在同一位创作者、同一批材料、同一段时长内，记录“材料事件”的总数与“被触动事件”的数量（后者可由创作者当场标记，并以停顿时长、注意力驻留、事后可复述性三项交叉验证）。应答理论预测两者应当高度同步——材料参与得多，创作也就推进得多；本文预测两者会显著分离，且**只有被触动事件的数量**与该阶段的方向改动次数相关。

## （三）与本栏接收端诸篇的分工：这一篇站在创作端

作者此前在本栏发表的若干篇——《工序的不可代偿性》《先占》《不可代理的残余》《无薪的劳作》《裂缝与在场》——处理的都是同一个方向的问题：**接收者**必须亲自跑完那段路，意义才发生，代劳不成立。本文与它们共用“不可代劳”这个大母题，因此必须把分工写明，否则读者有理由认为这不过是同一句话的第六次说法。

分工是**回路的两端**：那几篇追问的是作品交到人手上之后，接收者这一端有什么不能被替他做；本文追问的是作品被做出来之前，创作者这一端有什么不能被替他做。两端的机制并不对称：接收端的不可代劳源于意义只能在接收者自己的时间里长出来，损害表现为“这不是我的事”；创作端的不可代劳源于因果归属只能挂在亲手操作上，损害表现为“我不想进工作室”。若有人能证明这两种损害其

实由同一个变量驱动——例如证明创作端的动力衰减也可以靠改善接收条件来缓解——那么本文与那几篇就应当合并为一篇，本文的独立性也就不成立。

## 附论二·可裁决预测：本文可以在哪里被判错

---

除正文第七节已给出的五年追踪设计之外，本文再给三条更短周期、可独立实施的预测。任何一条被证否，本文的核心命题都需要实质性修改。

**预测一（因果归属的排他性）。**在实验条件下让两组创作者面对完全相同的一批"意外效果"：A 组由自己的手（通过一个引入随机扰动的画笔）造成，B 组由算法生成后呈现，两组在视觉上不可区分且不知情。本文预测：A 组的方向改动率显著高于 B 组，且这一差异在被告知真相后仍然保持（因为注册发生在预反思层）。若两组无差异，或差异在告知后消失，则"因果归属是触动的前提"这一核心机制被推翻。

**预测二（脱扣的节点特异性）。**把 AI 分别插入回路的四个节点——只做前期调研（不入回路）、代为执行操作、代为筛选意外、代为给出改向建议——追踪四组的创作持续周期。本文预测：只做调研的一组与亲力亲为组无显著差异，而代为执行与代为改向两组显著恶化，且恶化程度大于代为筛选组。若恶化程度与插入位置无关（即只要用了 AI 就一样坏），则本文的"节点脱扣"框架是多余的，一个笼统的"依赖假说"就够了。

**预测三（再踩的可逆性）。**对已出现动力衰减的创作者施加一次"强制亲手"干预（一段时间内只用不可撤销的材料工作，如湿壁画、直接雕刻、一次成像），本文预测动力指标可在数周内部分回升，且回升幅度与该期间的被触动事件数量相关。若强制亲手不带来任何回升，则说明衰减另有来源，本文把它归因于回路脱扣是错的。

## 参考文献

---

- Csikszentmihalyi, Mihaly (1990). Flow: The Psychology of Optimal Experience. Harper & Row.
- Flusser, Vilém (1985). Ins Universum der technischen Bilder. European Photography.
- Gallagher, Shaun (2005). How the Body Shapes the Mind. Oxford University Press.
- McLuhan, Marshall (1964). Understanding Media: The Extensions of Man. McGraw-Hill.
- Merleau-Ponty, Maurice (1945). Phénoménologie de la perception. Gallimard.
- Merleau-Ponty, Maurice (1961). L'Œil et l'Esprit. Gallimard.
- Rosenberg, Harold (1952). "The American Action Painters." ARTnews, December 1952.
- 石涛(清). 苦瓜和尚画语录（《画谱》），清康熙年间刊本。
- Schön, D. A. (1983). The Reflective Practitioner: How Professionals Think in Action. New York: Basic Books.  
（本轮所对质的最贴近对手：情境的对话）
- Ingold, T. (2013). Making: Anthropology, Archaeology, Art and Architecture. London: Routledge.
- Malafouris, L. (2013). How Things Shape the Mind: A Theory of Material Engagement. Cambridge, MA: MIT Press.
- Boden, M. A. (2004). The Creative Mind: Myths and Mechanisms (2nd ed.). London: Routledge.（计算创造力的经典分析，替代原稿中未能核实的两条引述）