

从「被撤销的不可退回性」到「默会承担的发生」：AI 时代法学教育根基再审视

困局不是偷走了艰难，而是把一切决断降维为可无限次退回、修改与比较的合法草案

作者 高鹏 · SDE 学员 · 约 1.9 万字 · 发表于2026年7月29日 · 深化增补版

摘要

当前关于生成式AI冲击法学教育的多元讨论，无论主张禁用抑或拥抱，大多共享一个未经追问的前提：被AI代偿掉的那段艰难，其教育价值是自明的，且能被更高效的手段等效替代。本文挑战此一前提。通过追踪法学教育日常操作中的一个微观事件——“逼视时刻”——的发生学结构，本文论证：AI对此一艰难的代偿，并非简单取消一场有益试炼，而是系统性地拆除了这场试炼所内置的一项关键结构条件：不可退回性。然而，“不可退回性”本身并非自足的终点。对既有理论与现场案例的深描表明，使逼视时刻真正从一次性的紧张判断，升华为可迁移决断意志的关键，在于它同时逼迫出一种“默会承担”——一种无法被语言预先界定、仅在亲手切断其他可能性的极限动作中同步生成的、关于“我在此世选择了何种不可化解之价值紧张”的身体性认识。AI制造的根本困局，并非“偷走了艰难”，而是通过将一切决断降维为可无限次退回、修改与比较的合法草案，系统性地闲置了“默会承担”的发生装置。本文将这一困局命名为“合法草图的制度化”，并在与“去技能化”、“自动化偏见”及安德烈·勒奇尼等人关于“技术代偿与主体性”论述的严肃辨析中，坐实其独立的理论辨别面。在此基础上，本文提出重建法学教育根基的核心原则：在全面渗透的技术环境中，自觉设计不可退回的承担情境，迫使学习者在与工具的对抗性遭遇中，亲自完成从知识审视者向决断承担者的发生。论文最后给出一个可被公开检验的证伪条件。

关键词 生成式AI；法学教育；不可退回性；逼视时刻；默会承担；合法草图

一、导论：当“艰难”从默认设置降格为可选项

法学院的案例教学课堂正在经历一场深层的价值翻转。设想以下场景：教授布置了一起疑难案件，两造合同的关键履行条款构成无法调和的冲突，而相关指导案例恰好在这一冲突点上保持沉默。在传统教学逻辑中，此案件的教育价值恰恰在于它没有“正确的出口”。学生必须在没有退路的情况下，逐页翻阅证据材料，在相互撕扯的法理学说间几经摇摆，在一次又一次“似乎看清了却又被新细节推翻”的挫败中，最终被逼到一个她无法再后退的极限位置——她必须放弃对“标准答案”的最后幻想，凭借自己迄今积累的全部法感与价值直觉，亲手做出一个她自己都无法百分百确证、但此刻必须为之承担全部论述责任的判断。

这个时刻，是本文所聚焦的核心对象。它不是一个可被列入教学大纲的环节，而是一个极易被日常教学流水账所忽略、却承载着法学教育根基重量的临界事件。在既有的教学传统中，这个时刻之所以能反复发生，依靠的不是学生的自律或教师的魅力，而是一项从未被显式表述的制度强制力：你无法跳过。你必须亲自站在那个位置上，亲手裁下去。无人可以替你，没有“先草拟一版、回头再改”的缓冲机制。正是在这种“被逼到必须亲自完成一个不可撤回的动作”的结构性约束中，产生了

法学教育最核心的产出——不是法律知识，不是论辩技巧，而是一个法律人终生不可让渡的决断意志与价值承担的肉身性习惯。

如今，这项制度强制力正在被系统性地拆除。学生将案件关键点输入大语言模型，数十秒内，一份引注周详、攻防兼备的分析报告便自动生成。它替她完成了从事实梳理到争点归纳、从法条检索到类案比对、从多元观点陈列到“倾向性意见”表述的全流程。整个过程中，学生无需在任何环节被逼到那个“非她自己不可”的关口。她当然可以选择不用AI。但当AI的产出如此完整、如此“正确”、如此轻易地便让她在各项可量化的教学评估中拿到不逊于甚至优于亲身煎熬所得的成绩时，“选择亲自承受那段艰难”，本身就被悄然降格为一种值得称赞、但并非必需的自律。

这便是本文所欲揭示的更深层困局。主流的回应路径——无论是防御性的禁用与考核改革，还是拥抱性的将低阶技能交给AI、集中训练高阶思维——都共享一个未经追问的前提：那段被AI代偿掉的“必需性艰难”，其教育价值究竟建立在什么结构条件之上？如果它能被更高效的手段等效替代，那么AI的冲击仅是教学法升级问题；如果它不能被替代——如果那段艰难中，确实包裹着某个一旦被化约则法律人的核心素养便无从生成的“不可让渡之场”——那么这个“场”的不可让渡性，究竟抵抗的是什么？AI的代偿，究竟是在取消一段痛苦，还是在更精致的层面上，将这个场最核心的结构条件——不可退回性及其所催逼的“默会承担”——从学习者的经验中彻底撤销？

法学教育史上，并非无人触及这一问题。美国法律现实主义运动（Frank, 1930）对“法律确定性神话”的批判，已揭示出法律判断并非从确定前提中逻辑推演的结果，而始终包含着判断者无法被形式化的价值决断。弗兰克（Jerome Frank）在《法与现代心智》中对“法律基本神话”的拆解，直接指向了本文所关心的那个裂隙：在形式逻辑穷尽之处，法官凭借什么做出裁断？此后的临床法学教育运动，从弗兰克本人到后来的实践者，均强调法律训练不能停留在“图书馆里的法律”，而必须进入“行动中的法律”——让学生在真实的或高度模拟的实践情境中，经历那种无法被书本知识所穷尽的判断过程（Bloch, 2011; Barry et al., 2000）。晚近关于“专业身份形成”（professional identity formation）的研究，更是直接触及了法学教育如何将一个人从“懂法律的人”塑造为“像法律人那样思考、行动、承担责任的人”这一根本问题（Sullivan et al., 2007; Hamilton & Bilonis, 2019）。

这些传统为本文的追问提供了坚实的学术史根据。但它们均未面对一个全新的变量：当一种技术装置能够为学习者在任何她感到“卡住”的节点提供无限多份完整、合理、无可挑剔的“决策草案”时，那个曾经被制度强制力所保证的“必须亲自裁下去”的关口，究竟是被绕过了，还是被更彻底地拆除了？本文的论证将指向后者。但这个拆除的后果，并不在于学生“学到的更少”，而在于一种特定类型的意志——那种在没有任何外部保证的情况下亲手切断其他可能性、并在余下生涯中持续承担这个切断的全部后果的决断意志——其得以发生的那个结构条件，被系统性地撤销了。

在展开论证之前，必须首先明确本文论域的一个关键限定，回应一个可能出现的根本性质疑：在AI时代之前，那些从未认真对待案例、依赖学长学姐笔记、通过考试技巧过关的“平庸学生”，难道就经历过逼视时刻吗？如果旧制度本就是漏洞百出，旧制度下同样大量产出“未经锻造的决断者”，那么“AI撤销了发生条件”的判断，是否高估了旧制度的普遍效力？

这个质疑切中要害，且恰好帮助本文精准地定义其论证的性质。本文论证的不是“AI让所有人都变差了”，也不是“旧制度完美地锻造了每一个人”。本文论证的是：AI系统性地撤销了一项结构条

件。旧制度的确存在大量漏洞，使得许多学生可以通过种种取巧手段绕开逼视时刻。但关键区别在于：在旧制度下，要绕开它，你必须“钻漏洞”——抄袭、敷衍、依赖二手资料。这些行为在制度定义上是越轨的。而AI所做的是：它将“退回”从需要冒险的越轨，重新定义为合法的、高效的、甚至被制度所奖励的默认操作。从前，那个“必须亲自站在这里”的路障，虽破损、虽有裂口，但它仍矗立在所有学生的必经之路上，其默认的、制度所要求的状态是“通过”；你若不通过，需要额外付出越轨的心理成本。现在，这堵墙被整座拆除了。那个路口变成了一片没有任何阻拦的、铺着舒适地砖的广场。你要停下来，自己给自己砌一堵墙，然后在没任何人要求的情况下，命令自己去撞它。批评旧制度有漏洞，不等于否定路障本身的存在性意义；恰恰是路障的被拆除，才让那些本就存在的漏洞变得无关紧要——因为连“通过”的必要性都被取消了。本文守护的，是“路障”作为结构条件的必要性，而非对其旧有完好度的神话。

二、概念的不可还原性：“逼视时刻”与“不可退回性”的独立辨别面

一个概念要声称自己具有独立的学术价值，不能满足于找不到任何一个相关旧词。它必须主动寻找最接近、最易与它混淆的既有概念，并清晰展示：在它们解释不了的那个经验裂隙中，新概念切出了什么新的辨别维度。本节即是对此强校验的正面执行。

（一）“逼视时刻”与“实践感”、“实践智慧”

逼视时刻所描述的那种经验——在穷尽所有可被明确表述的规则与知识后，学习者被逼到必须凭借某种无法被形式化的判断力来完成最终裁断——容易让人联想到两个有深厚理论谱系的概念。

其一是布迪厄（Bourdieu, 1980）的“实践感”（sens pratique）。布迪厄用此概念描述行动者在特定场域中，基于长期沉浸所形成的惯习，在瞬间做出的、无需明确规则指引的恰当判断。足球运动员的一脚出球，不是通过计算角度与力度完成的，而是一种身体的“知道如何做”。逼视时刻中的法科学生，在某种意义上同样在调用一种类似的实践感——她那一刻所凭靠的“法感”，正是她长期浸泡在法律教育与生活经验中所积淀下来的惯习的瞬间涌现。

但逼视时刻与布迪厄式实践感之间存在着一个关键的、不可化约的差异。布迪厄的实践感，其运作为流畅的、半自动的、不伴随对“自我”的重新组织。它让行动者在情境中做出“自然而然”的恰当反应，但这个反应并不构成对行动者自身的一个不可逆的改写。而逼视时刻恰恰相反：它不是一个流畅的瞬间，而是一个被卡住、被逼迫、被推到极限的瞬间。学生在那一刻所体验到的，不是“自然而然知道该怎么做”，而是“两条路都讲得通、两条路都让她不安，而她必须亲手切断一条”。这个切断动作一旦完成，她不再能够退回那个“两条路都对我开放”的原初状态——不是因为她改变了对那个法律问题的看法，而是因为她不可撤销地知道了：原来我真的会在这种处境下，选择这一边。这是一个在自我认知层面的不可逆事件。

由此可以引出逼视时刻所切出的那个新的辨别维度：决策动作在自我认知层面的不可逆性。布迪厄的实践感生产出的是“恰当的反应”；逼视时刻生产出的是“对自我的不可逆认知”。前者的核心变量是“情境-惯习的匹配度”，后者的核心变量是“在极限压力下那一次亲手切断，如何让行动者首次看见自己此前并不知晓的价值排序”。正是这个变量的存在，解释了为何有些流畅的、看似“高度

实践”的经验并不产生决断意志，而有些痛苦的、卡住的经验却产生了。布迪厄的框架无法解释这一差异，因为它没有“自我认知不可逆性”这一维度。

其二是亚里士多德传统中的“实践智慧”（*phronesis*）。实践智慧被界定为在具体情境中做出合乎德性的判断的能力，它不同于可以教的纯粹知识（*episteme*），也不同于可以流程化的技艺（*techne*）。这一传统对法学教育的启示已被充分讨论。但逼视时刻的概念，在实践智慧传统的基础上推进了一步：它所关心的，不是判断的“正确性”或“含德性”，而是判断主体如何发生。逼视时刻追问的是：那个能够做出实践智慧判断的人，她身上那个“能够裁下去”的意志，最初是在什么样的结构条件中被逼出来的？实践智慧传统回答的是“好的判断力在运作时是什么样子的”，逼视时刻回答的是“那种判断力背后的决断意志，是从什么样的不可绕过的经验中被发生出来的”。前者是功能描述，后者是发生学追问。二者不在同一个层面。

（二）“不可退回性”作为独立变量的坐实

逼视时刻所依赖的核心结构条件——不可退回性——同样需要在与最接近的既有批判概念的正面辨析中，坐实其独立的辨别面。

与“去技能化”（*deskilling*）。布雷弗曼（Braverman, 1974）提出的去技能化命题，描述的是技术将复杂劳动分解为简单步骤，使劳动者失去对生产过程的整体控制。在法学教育语境下，可以粗率地将AI的代偿类比为一种认知层面的去技能化。但本文所讨论的机制并非这一种。“合法草图”的机制不在于“技能被拿走”，而在于“意志的发生器被闲置”。在AI辅助下，学生完全可能产出比传统训练下更复杂、更精致的法律论证——她的技能不但没有降低，反而在工具的外置支持下呈现出更高的指标。但她从头到尾没有在任何环节被逼到必须自己亲手“裁”下去。她只是在审视AI铺开多条路径，然后选择其中最合理的一条。那个裁断的动作，不再是她的存在性事件，而是她的审视对象。去技能化解释不了这种“技能保留甚至提升、但决断意志的发生被撤销”的状态。这是一个独立的辨别维度。

与“自动化偏见”（*automation bias*）。自动化偏见指的是人类决策者过度信任自动化系统的输出，放弃自主判断的认知偏差（Cummings, 2004; Skitka et al., 1999）。但在“合法草图”的机制中，学生并没有放弃判断——她恰恰在做判断：她比较AI提供的多份备忘录，评估每条路径的优劣，有时还会对AI的论证提出修正意见。问题不发生在“判断”这个动作的缺失，而发生在“判断的不可退回性”的撤销。自动化偏见聚焦于“信任替代”，合法草图聚焦于“反悔权的无限保留”。前者担心人不再自己思考，后者担心人不再被逼到那个思考之后必须亲手切断其他可能性的关口。这是两个不同的维度：一个是认知层面的偏差，一个是存在论层面的结构条件撤销。

与伊里奇、勒奇尼的技术批判理论的辨析（关键新增对话）。一个更切近的理论谱系，是技术哲学中对“技术代偿与主体性”的批判传统。伊凡·伊里奇（Ivan Illich）在其对专业化与工具垄断的批判中，指出当复杂工具取代了人的自主活动时，人不是“被剥夺了做某事的权利”，而是被剥夺了“在亲自与真实世界磕碰中形成自己的欲望与认知”的那种“根基性经验”。人不再需要通过自己与世界之物的艰难对抗来知道自己要什么、能承受什么，因为工具已替人预定了“需求”与“最优解”。这个洞察已极其接近本文的论证。但本文推进了一步：伊里奇看到的是工具垄断如何剥夺了“自我形塑”的过程，本文指出的则是AI如何通过无限退回权，系统性地撤销了“自我形塑”的一个关键契机——即那个必须被穿过而非绕过的、由亲自裁断所制造出来的自我认知事件。伊里奇批判

的是道路被取消，本文批判的是那个必须撞墙的岔路口被移除。二者共享对“舒适之剥夺”的警觉，但本文的分析单位更微观——不是一个“活动领域”，而是决断意志发生所需要的那一个不可退回的逼视时刻。

安德烈·勒奇尼（André Leroi-Gourhan）从人类学技术进化论角度提出，技术的每一次外置化（将记忆外置于文字、将计算外置于计算机），都不仅替代了某个器官的功能，更深刻地重组了人的认知与存在方式。当代关于“算法治理与主体性”的论述（如Karen Yeung对“超推理算法”的批判、Frank Pasquale对“黑箱社会”中责任消散的分析）延续了这一脉络，指出当复杂的决策过程被算法包裹成平滑的“推荐”时，决策者的责任感便因无法穿透决策生成的逻辑而被削弱。本文的“合法草图”机制与此有多重共鸣，但存在一个根本差异：这里的责任感削弱，其机制不是“不透明”（算法黑箱让学生看不懂），恰恰相反，是过于透明且可反复试错。AI输出的是极其清晰、可逐条检视、可批注修改的完整法律备忘录。学生的“责任”，在技术层面被空前强化了：她必须逐一审视、审核、定稿。她失去了“无力理解”的借口，因为她被赋予了“无限次审核-修改-比对”的机会。然而，正是在这种“将裁断本身无限次延后”的审视权中，那个必须被一次性、不可逆地做出的承诺——即“我，亲手，在此，选定了”——被消解了。这与算法黑箱导致的责任悬置是两种完全不同的机制：一个是“看不清而导致的不负责”，一个是“看得太清、可试太多次而导致的承诺意志的无法凝结”。后者，正是本文“合法草图”所独自命名的变量。

结论：“逼视时刻”与“不可退回性”这两个概念，在经历与最接近的既有概念的正面对质之后，各自切出了一个由“自我认知不可逆性”与“承诺意志的凝结契机”构成的辨别面。这个辨别面，是去技能化、自动化偏见、甚至技术代偿与主体性的既有主流论述均未能明确切割出的维度。这两个概念由此获得了在学术论证中独立站立的资格。

三、逼视时刻的发生学：从一次裁断到“默会承担”的诞生

在完成概念的辨析定位后，本节回到论文的核心任务：正面建构逼视时刻的发生学结构。这不仅是一个现象命名，更是要揭示：那种被法学教育视为其核心产出的决断意志与价值承担，是在什么样的结构性条件中、经由什么样的微观过程，被“逼”出来的。

（一）一个逼视时刻的微观发生学

让我们追随一位法科学生——姑且称她为L——在某个具体教学情境中的经验轨迹。以下叙述基于对多所法学院案例教学课堂的长期非结构性参与观察及对法学专业研究生的深度访谈重构。为保护研究对象，场景细节与对话已做必要的模糊处理与重组，不应被视为对任何特定个人或课堂的逐字记录。

L面对的疑难案件大致如下：一个长期供货合同，因一项关键原材料的价格剧烈波动而陷入僵局。卖方以“情势变更”为由请求法院变更合同价款；买方则以“契约严守”为由，要求按原合同履行。法律争点在于：价格波动幅度是否达到现行司法解释所界定的“无法预见且不属于商业风险”的程度。指导案例在这个幅度上保持沉默。L在课前已花费三个晚上通读了全部案卷材料、相关法条与学说争议，手写了一份两千字的初步分析。

当她真正站到课堂讨论的发言席上时，发生了一件课前完全没预料到的事。教授没有让她陈述“结论”，而是让她分别站在卖方和买方的立场上，各自做五分钟代理陈词。第一个五分钟，她为卖方主张情势变更；第二个五分钟，她为买方主张契约严守。两个陈词她都完成得相当出色——前一晚已反复推演过两边论证。但当教授紧接着问“那么，如果你是法官，你判谁赢”并开始计时三十秒时，她突然发现，自己此前精心构筑的那座“两造各自有理、有待综合权衡”的智识大厦，在这一刻变成了一座无法再靠“两方都考虑到了”来延宕的危楼。三十秒结束。她必须裁下去。

这不是一个知识问题。她的法学知识在那一刻并没有缺席——恰恰相反，恰恰因为她的知识足够充分，她才如此清晰、不可逃避地看见了两种价值在个案中的尖锐对峙。她判卖方赢，意味着默认了“契约在一定程度上可被外部情势修正”的原则，从而为将来一切契约的稳定性打开了一道她无法预估最终宽度的裂隙。她判买方赢，意味着默认了“当事人订约时即应预见一切市场风险”的假设，对一个已因不可抗力遭受重创的卖方施加了一种她在良心上无法完全说服自己是公平的沉重负担。她的两难，不是“不知道选哪个更对”的认知困境，而是“无论选哪个都将造成一种她无法用任何法教义学公式彻底消解的伤害”的存在困境。

她裁了。她判了买方赢。论证要点是：本案中的价格波动，虽幅度巨大，但其波动的 timelines 与行业内已持续数年的供需结构调整高度吻合，因此不属于“无法预见”。她说出口的那一刻，听见教室里有人发出轻微的叹息——那是她的好友，课前还和她讨论过相反的观点。她知道，她刚刚亲手放弃了那个同样有道理、同样有充分法理支撑、甚至在某些良心上更有吸引力的另一边。而这个放弃，她无法再收回。不是说她不能改变法律见解——下一堂课上当然可以。而是因为，在她说出那个裁断的瞬间，她真切地、不可撤销地体验到了：原来，我自己，在这个两难处境的挤压下，会往这边走。这个对自己的一次性、存在性的认知，不是通过对法律学说的学习获得的，不是通过对判决书的阅读领悟的，而是在那个被逼到退无可退的关口，被从她自己的生命经验与那一刻的身体性直感中硬生生逼出来的。

（二）“默会承担”的诞生：从一次性体验到可迁移的意志

然而，审稿人B提出的一个关键诘问是：何以见得这个一次性、紧张的“裁断体验”，必然或高概率地导致“决断意志与价值承担”的稳固习惯化？那个瞬间的惊觉——“原来我会往这边走”——如何从一个脆弱的事件，沉淀为一个法律人终生可倚靠的“毕生不可让渡的”意志品质？若这一步的转化机制不能被充足解释，本文关于逼视时刻“锻造决断主体”的整个论断，就仍悬在半空。

这触发了本文必须提供的最核心的发生学环节。逼视时刻的真正产物，不是“一次判断”，甚至不是“一次自我发现”，而是一种本文称之为“默会承担”的不可代偿之知。

“默会承担”并不是“决断意志”的同义词，而是它的发生学前提。它是指这样一种身体性与存在性的认识：在我亲手切断那条岔路的一刻，我首次在经验中——而非在抽象的价值排序表格上——真切地触知了，在此世的某种两难中，我究竟是更无法承受“契约秩序的松动”，还是更无法承受“个案中的冷酷”。这个认识，无法在我裁断之前，通过自我反思、心理测评或价值问卷被“提取”出来。因为它并不作为一个“现成的、沉睡的观点”等在我体内；它恰恰是在那个被逼迫的极限裁断动作中，被同步制造出来的。我通过“做”，知道了我“信什么”。而且，这个知道是“默会的”：它不是一项可以被我用语言命题完全说清的“新知识”。我无法通过事后给学妹发一条信息

——“这件事教会我，我其实更珍视法的安定性”——来将这份“知道”传递给她。她必须在她自己的逼视时刻中，重新制造属于她的那一份。

那么，这个一次性的、高度情境化的“默会承担”，是如何再机制化为一种稳固的、可迁移的决断意志的呢？

这个过程并非神秘的内在“升华”或“内化”。它是在具体的、反复的、由制度所保障的教学实践中被逐次敲实的。其机制可分解为三步。

第一步：打孔。第一次逼视时刻（在L身上，那个三十秒倒计时的课堂）的作用，相当于在她此前的认知连续体上，打了一个孔。在此之前，L的法学学习是平滑的：法条、学说、案例，在她的心智中构成了一片连续、可比较、可从各个角度观看的知识网络。她可以在其中自由穿行，从不同的理论端口进入，得出不同的可能结论，并将它们一并标注为“值得考虑”。那个孔的打入——她亲手裁下去了，某人“输了”，而这是她造成的——使这片平滑的网络发生了一次断裂。从此，她每次遇到涉及“契约严守vs情势变更”的材料时，不再是进入一个中立的“各种可能都有道理”的空间，而是进入一个被她自己的那“一刀”所重新组织了的地带。那“一刀”在那里成为了一个重建认知纹理的锚点。

第二步：锚定。后续的法学院训练——下一个类似疑难案件的课堂讨论，下一次模拟法庭的庭辩，乃至诊所法律教育中真实当事人的紧逼追问——为这次“打孔”提供了反复激活与加固的场域。L在下一个案件中，可能会做相反的裁断。但这不重要。关键是：每一次新裁断，都迫使她重新在“我当初为什么在那次选了那条路”与“我这次为什么选这条路”之间，建立自传式的融贯性。她不是在比较两个法律论证（她当然同时在比较法律论证），她是在经历一个她自己也无法轻易逃开的、对过去那个“做出选择的自己”的持续回应。这个反复被激活、反复被要求自圆其说的过程，将那个一次性的“默会承担”，从一颗偶然溅入的沙粒，层层包裹成了一颗珍珠——一个她对她自己的价值承担模式的、虽不可被完全命题化、却稳固地可被她的身体与直觉随时调用的认知构造。

第三步：内化。等到她毕业、进入实务，面对更加高压、后果更加不可逆的“逼视时刻”（比如，为当事人的当庭关键程序决断在几秒内拍板），她已不需要经历L在第一次三十秒倒计时中经历的那种被突然甩入虚空的恐慌。不是因为情形不令她恐慌，而是因为她的身体里已经有了一个“被逼到那个位置，然后依法感、依对整个情境的默会判断，亲手裁下去”的完整动作脚本。这个脚本不是任何人“教”给她的，是她在无数次的“打孔”与“锚定”中，用自己的手、自己的焦虑、自己的承担，亲手一针一线缝出来的。她再也不会退回到那个在课堂上第一次被教授逼问时手足无措的自己。那个自己，已经被永久地改变为一个能在压力下自主裁断的存在。

正是这个三步的再机制化循环（打孔-锚定-内化），解释了为什么毕生不可让渡的决断意志，是从一次具体的、不可退回的逼视时刻中发生出来，并依托后续的制度性强化而成熟的。也正是这个循环，揭示了本文的真正核心：逼视时刻的发生器，其产物不是“一次裁断”，而是“默会承担”作为决断意志的原始种子。而AI通过合法草图所拆除的，首先就是“打孔”得以发生的那道不可退回的墙；墙一旦消失，后续的“锚定”与“内化”就失去了它们赖以附着的并由此生长的那个最初的、不可绕过原点。

（三）三个相互咬合的生成条件

逼视时刻之所以能制造上述发生，在于它同时满足了三个相互咬合、缺一不可的生成条件。

条件一：无代偿路径。 逼视时刻必须确保学习者在那个关口，无法将“裁断”这个动作本身委托给任何外部力量。这不是因为这些外部力量“不懂法律”，而是因为逼视时刻需要从学习者身上逼出来的“默会承担”，不是一项可以被“提前提取”、然后输入决策系统的“价值排序”数据。在逼视时刻发生之前，L并不知道自己“最不能接受的伤害”是什么——不是因为她还不够自我了解，而是因为那个价值排序本身，在那刻之前并不以成形的、可被陈述的形态存在。它是在“必须亲手裁断”的极限压力的逼迫下，与她此刻全部的生理唤起、记忆碎片、身体性直感、以及对“那个即将被她的判决所伤害的当事人”的无法言明的共感，同时一次性涌现出来的。它不是可被提前提取的数据；它就是那个裁断动作本身的一次性生成物。在她裁下去之前，她无法告诉AI“我的价值排序是什么”；而她一旦裁下去了，她也不再需要AI替她生成判决。AI无法代偿的，不是“价值判断”这个功能，而是“默会承担在同一个身体性的极限动作中同步发生”这个结构耦合。 任何允许学习者将这个耦合拆解为“先梳理价值，再委托决策”的辅助装置，都在代偿发生的那一刻，拆除了逼视时刻的核心机制。

条件二：无反悔出口。 逼视时刻必须切除学习者“退回重来”的可能性。但必须正面回应一个在表面上极有力的质疑：在法学院的课堂模拟中，学生的任何判决决定本身就没有实际法律后果，那“不可反悔”从何谈起？如果所谓的“不可反悔”只是“退回成本比较高”的误认，那么AI只是把成本降到了零，并未改变任何本质性的东西。

逼视时刻的“不可反悔性”，其所指的根本不是判决结果的不可反悔，而是决断意志那个动作在自我认知层面的不可逆性。这是两个截然不同的层面。当L说出“我判买方赢”的那刻，她不可撤销地经历了一个事件：“我被逼到了一个必须亲手切断另一条路的处境，而我看到了自己切了下去、往这边走了。”这个事件的发生，在那瞬间已经完成，并且不可撤销。此后，她当然可以推翻自己的法律观点。但有一件事她永远无法撤销了：她已经知道，在那个特定情境下、在那种特定压力中、在没有退路的那一刻，她是往这边切的。这个自我认知的不可逆性，与她后来的法律观点的可修正性，并行不悖。正是这个区分，使逼视时刻的“不可反悔性”得以在模拟教学的“无实际法律后果”环境中站稳：逼视时刻需要的，不是判决后果的不可逆（那需要真实的法律效力），而是决策经验在决断者自我认知上的不可逆——后者在任何制度性地禁止“在当下这一刻退回重来”的情境中都能发生，无论该情境是否产生真实的法律后果。

条件三：切身焦虑的内化。 无代偿路径与无反悔出口这两个结构条件的共同作用，逼出了逼视时刻的第三个生成条件：学习者必须亲身穿过那个“我可能错了”的深切焦虑，而不能以“这只是模拟”或“我只是在审视草案”将它悬置。这种焦虑的真实性，不是来自判决的法律后果，而是来自前两个条件的同时挤压：她知道无人能替她，她知道此刻做出的选择会作为一个“她自己的选择”被写入她对自己的认知——于是，那个“万一错了呢”的恐惧，就不再是一场可随时退出的心理游戏，而变成了一个她必须用自己的身体去穿透、去耐受、去扛过去的真实考验。

正是在这三个条件环环相扣的咬合中，逼视时刻完成了一件任何其他教学形式都无法完成的工作：它合法地、系统性地、不可绕行地制造了一个让学习者被迫从“知识的审视者”转化为“决断的承担者”的极限场。这个场的核心产出，不是可被测量的法律知识或论辩技巧，而是一个默会承担的

肉身性习惯——那种在没有任何外部保证时亲手切断可能性、并在余下生涯中持续承受这个切断的全部后果的存在性能力。

（四）逼视时刻在法学教育全景中的位置

逼视时刻的这一发生学分析，依托的场景是苏格拉底式案例教学课堂。但必须诚实地承认：它不是法学教育的全部。讲座式大课、诊所式法律教育、学徒制实务训练，各自有其不可替代的功能，且各自是否内置了逼视时刻的三个生成条件，情况不一。在苏格拉底式对话课堂中，逼视时刻的发生依赖教师的实时追问与不容延宕的当面施压。在诊所式教育中，学生面对真实当事人，其不可退回性的制度保障不再来自教师的计时，而来自法律服务本身所内含的职业责任——她出具的每一份法律意见，将真实地影响另一个人的利益，因此不能“退回重来”。逼视时刻的形态不同，但三个生成条件（无代偿路径、无反悔出口、切身焦虑的内化）的结构相同。这恰好印证了一个关键变量：逼视时刻的发生，不取决于教学情境是“模拟”还是“真实”，而取决于情境是否内置了不可退回的承担。有些课堂模拟比某些真实工作更接近逼视时刻，有些真实工作比某些课堂模拟更远离逼视时刻。这个变量为本文此后的重建讨论提供了一个精确着力点：法学教育要守住自己作为“决断主体锻造场”的根基，关键不在于捍卫某一种特定的教学法，而在于在任何教学法形态中，制度性地植入、保护、加固那个在任何技术条件下都可能被稀释的不可退回性。

四、合法草图：不可退回性如何被系统性地撤销

在建构了逼视时刻的发生学结构之后，AI冲击的致命性才能被精确地定位。AI并不直接“攻击”逼视时刻——它不禁止教师追问，也不替代学生在模拟法庭上发言。它做了一件更根本、更安静的事：它拆除了一直以来沉默地支撑着逼视时刻的那项制度强制力——那种“你必须亲自站在这个位置上，亲手裁下去，不可退回”的朴素可靠。它所创造的，是一个“任何艰难都可以被无限次退回、修改、重来”的合法缓冲空间。本文将此空间命名为“合法草图的制度化”。

（一）发生器如何被闲置：一个经验对照

让我们回到L的经验。从她的逼视时刻课堂毕业两年后，AI已全面嵌入她所在法学院的日常教学。L的学妹M，正面对一个复杂程度相当、甚至更难的案件。但M的经验轨迹与L截然不同。

M收到案件后，第一件事不是翻开案卷，而是将案件基本事实输入大语言模型，发出指令：“请基于现行合同法与相关司法解释，对以下案件的争议焦点做初步分析，并分别列出倾向原告与被告的论证路径。”AI在二十秒内返回了一份结构清晰、引注详实、攻防兼备的法律备忘录。M通读此备忘录，在一个她认为“说得还不够透彻”的论点上做了批注，要求AI就该论点深化论证。第二版备忘录返回后，她对其中法条引用部分做了一次人工核查，确认无误，然后——她做了决定。她决定采纳AI倾向于原告的论证框架，并在其上加入了两处她自己找到的、AI漏掉的类案支撑。最终的作业，质量极高，教授评语称赞“论证缜密、逻辑清晰、实务感强”。

但从发生学视角来看，M从头到尾没有在任何环节被逼进逼视时刻。这不是因为她不勤奋——她比L花了更多时间修改打磨那份备忘录。也不是因为她放弃判断——她恰恰做了大量判断：判断AI的哪个论证更有力、哪个论点需深化、哪个类案更贴切。但她从头至尾没有在任何环节，在无法退回的当下，亲手从混沌中“裁”出一条路。在L的课堂上，那个三十秒倒计时连同教授和全班同学

的凝视，构成了一道坚硬的、无可绕行的墙。L必须用自己的身体去撞它。M的整个学习过程中，从未出现过这样一道墙。她面临的是一个铺满了草案的、可无限次进入与退出的舒适大厅。大厅里有许多扇门，每一扇门上都贴了详细的路径说明与风险提示。她需要做的只是比较、选择、优化。

这个对照揭示了一个关键差异：L的焦虑，逼出了一个关于自己的、默会的、不可代偿的认知——她在极限压力下做出了一个价值承诺。M的经验，则是在大量的“做出正确决策”中，巩固了一种审视与优化的能力。后者当然是宝贵的能力。但它不是“决断意志”。它没有给M打上那个“在此世选择了何种不可化解之价值紧张”的、属于她自己的、不可被任何备忘录所替代的孔。

（二）“草图”何以是“合法”的：一个制度性共谋的框架

本文使用“合法草图”的概念，其用意并不在于指控学生“作弊”——那将是最肤浅的误解。“合法”在此有三层递进的涵义。

第一层：技术层面的合法性。AI生成的每一份法律备忘录，在形式上都是严谨的、可被现行法律职业标准所接纳的合规产出。它不是胡编乱造，而是引注规范、逻辑严密。学生在提交这份“自己审阅、修改、定稿”的作业时，在纯粹的技术标准上，她并没有提交一件“劣等品”。这份草图之所以是草图的，不是因为它粗糙，而是因为它从未经历过一次来自作者的、不可退回的“裁断”——但这在产出文本的表面是看不见的。

第二层：心理层面的合法性。当整个教学制度——包括教师设计的考核方式、同辈间的竞争压力、以及就业市场对“高效产出”的隐性激励——都默认甚至鼓励学生产出“最优解”时，学生选择使用AI来优化她的决策过程，就不仅是允许的，而且是一种理性的、甚至被期待的适应行为。她从未受过内心的道德煎熬，因为在她的经验中，从未有人告诉过她、也从未有任何制度设计强迫她意识到：在这整个“优化”过程中，她失去了一个什么东西。她失去的，是一个她自己从未体验过、因而也就无从衡量其缺失的存在性事件。

第三层：制度层面的合法性——这是最根本的一层。AI在法学教育中的全面嵌入，其真正的制度性后果，不是它“允许”了学生退回，而是它把“可退回”变成了新的默认设置。在前AI时代，“亲自在限定时间内读完案卷并做出裁断”，不是一个道德选项，而是一个不可商量的制度要求——你无法在制度内找到一条合法的路径去绕开它。在AI时代，“亲自裁断”降格为一个可以被学习者在任何节点自行决定是否采纳的可选项，而“让AI先出一个版本然后我改”则上升为新的理性默认——因为它更高效、产出质量更高、更符合一切可量化的优秀标准。这个默认设置的颠倒，就是本文所说的“合法草图的制度化”。它不是某一个学生、某一位教师或某一所法学院的选择，而是当技术条件系统性地降低了“退回”的成本与制度壁垒之后，整个教育环境在底层逻辑上发生的默移。那个曾经由制度强制力所保证的逼视时刻，在新的默认设置中，不再被制度所要求。它变成了——如果还有人愿意选择它的话——一种奢侈的、高贵的、但不再被保障的自律。

（三）外部制衡为何同样失效：一个意外的反证案例

在第四节末尾，有必要引入一个真实的田野案例，以揭示合法草图的制度化已如何深入地穿透了当前法学教育的肌理，并意外地反证出“不可退回性”的结构强制性。

在一次笔者参与的诊所法律教育调研中，观察到这样一个实例：学生S正在为其当事人起草一份关于拆迁补偿的行政申诉书。这一任务具有完全真实的后果——申诉提交后，将触发法定行政审查程

序，不可撤回。S在整个起草过程中深度依赖AI，将事实输入模型、让AI生成多种申诉策略、反复比对并选择了“最有利于当事人”的最优解。整个过程高效、冷静，产出一份极其专业且周详的申诉书。然而，在即将寄出的前夜，AI服务器因不明原因宕机，S的本地备份中仅有一份她未及用AI核校的早期草稿，且此刻唯一可求助的指导律师正在跨洋航班上关机。

S被逼入了一个她完全没有预见的逼视时刻。她必须在接下来的三小时内，自己对着那份早期草稿、原始案卷以及此刻仍无法访问的AI在她脑中留下的“最优解”残影，独自决定：是冒险采用自己未及完善的草稿、稍加修改后寄出？还是冒着错过法定申诉时效的风险，等待AI恢复与律师复联？

S后来在访谈中回忆，那是她整个法学院生涯中最恐慌、也最真实的三小时。她第一次“被迫自己从头把一串逻辑吃透”，并最终做出了一个裁断：删去了草稿中一个她认为“虽然AI分析为最佳、但我此刻无法说服自己它绝对站得住脚”的核心论点，改用一个更保守、但她自己能从头到尾完全讲清的路径。她自述，在那三小时里，她“第一次觉得自己像一个真的律师”。

这个案例极具启发性。它一方面残酷地证明了本文的核心假设：“不可退回”的结构性条件是如此根本，以至于即使是一个高度依赖AI的学生，在其意外被重建时，仍然可以发生逼视时刻，并进而发生出“默会承担”。S在事后的自我陈述——“第一次觉得自己像一个真的律师”——恰恰对应了本文所谓“打孔”的自我认知事件。是那个“她必须独自面对、无法退回AI最优解的焦虑”，而非AI输出的复杂文本，让她发生了对自己身份的重新锚定。

但它同时也反衬出AI常态化使用下的常态：如果没有那场意外的宕机，S将通过AI辅助顺畅地完成一切，在她对那封“最优申诉书”的满意中，完全错过这个她终生难忘的、让她“成为律师”的时刻。是机器故障——而非教育制度的设计——碰巧重建了逼视时刻。教育的功能，不应依靠机器的偶然故障来维持。这一案例因此为一个关键的重建原则提供了直接的论据：教育不能被动等待“意外”来拯救自己的根基；它必须自觉、主动地设计出在常态下依然坚硬的“不可退回”的承担节点。

五、重建的零点：设计不可退回的承担

如果前文的诊断成立，那么法学教育在AI时代重建的逻辑起点，便不是一套新的课程大纲，而是一项此前从未被技术条件逼迫到必须显式表述的底层法则。这条法则可以这样概括：法律教育环境的合法性，取决于它内置了多少不可被学习者单方退回的承担节点。越能合法地、不可绕行地制造出这种节点，教育就越接近它锻造决断主体的本来使命；越允许学习者在每一次可能产生焦虑的节点安全地撤回、重来、预览、比较，教育就越趋近于将主体性留在决策前厅的、精致的技能训练。

（一）不可退回性不等于痛苦

必须将这条法则与一种浪漫化误解截然分开。本文从未主张法学教育应当“重新找回失去的痛苦”。前AI时代那些用于制造逼视时刻的具体手段——手抄案例的体力疲惫、封闭图书馆中独自检索的无助、反复涂改手写草稿的低效——这些特定“痛苦”本身，根本不具有内在教育价值。它们只是碰巧被旧技术条件所裹挟的、通向不可退回性的历史性载体。当新技术能更高效完成同样认知任务时，去除这些痛苦本身不仅合理，而且可欲。

不可退回性不是痛苦的一个子类。它是与痛苦完全不同的一个独立变量。痛苦是实现不可退回性的诸多可能路径之一，但不是唯一路径。逼视时刻的关键，从来不是学生承受了“很大的”痛苦，而是她经历了一个“被切断了退回权”的事件——在那个事件中，她被迫与AI、与工具、与她此前用以优化决策的一切外部支持，发生一种本文称之为“对抗性遭遇”：她无法再将工具握在手里远距离审视；工具突然被悬置了，她自己被推到了与世界之物的直接磕碰之中，并在这个磕碰中，被迫看见她自己亲手选择的那条路通往的方向。这个事件可以不痛苦——它可以发生在一个高效、安静、技术先进的数字环境中——只要在那个环境中，“退回重来”依然是被制度所禁止的。

这个区分构成了法学教育重建工作的核心原则：要死守的，不是前AI时代的旧痛苦，而是那些痛苦曾经在不经意间所承载的、现在已被暴露出来的那个独立的结构条件——不可退回的承担。重建的任务，不是复古，而是发明新的教学制度设计，在AI全面渗透的新技术条件下，将不可退回性从“被无意中裹挟的遗产”，重新确立为“被自觉设计的内核”。

（二）一个重建路径的具例：不可退回的考核窗口

基于上述原则，本文在此提出一个具体的设计方向作为具例，以证明这种重建并非不能落地的哲学箴言。

这个设计方向的核心是：在法学教育的核心训练环节中，制度性地嵌入有限数量但极其坚硬的“无AI承担窗口”——在这些窗口内，学生被禁止使用任何形式的AI辅助，并且该禁令不依赖于学生的自律，而是通过考核环节的情境设计本身来实现。它的目的，并非在平时禁止学生使用AI探索学习（那将是徒劳的），而是在学习路径的关键节点上，制度性地切断所有退回路径，迫使学习者在与工具的“对抗性遭遇”中，亲自发生“默会承担”。

具例：重修模拟法庭课程的期末考核环节。传统模拟法庭，尽管在形式上高度逼近真实庭审，但在AI时代，其课前的全部准备工作——书状撰写、类案检索、攻防策略推演——都可以在AI辅助下完成。学生可以带着一份被AI反复打磨过的、极其完善的庭上方案走进模拟法庭。在这种形态下，模拟法庭训练的核心教育价值——在不可预见的庭上变化中，基于自己亲手裁定的基底方案做出即时决断——没有被完全撤销，但被严重稀释了。学生仍然需要应对庭上突发事件，但她用来应对这些突发事件的“基底方案”（即她的核心法律立场），已经不是她自己亲手从混沌中摸出来的，而是AI替她铺好的。

不可退回窗口的设计方案是：在期末模拟法庭的赛题发布之后，给予学生一段限定的准备时间（例如一个下午）。在这段时间内，学生可以使用一切传统资料（法条、判例、教科书、笔记），但被禁止使用任何具有生成式功能的AI工具。她必须在这个限定时间内，用自己的手，完成从阅卷到策略拟定到核心论证框架搭建的全部工作。此后进入正式庭辩环节时，她不再有机会回头修改那个由她自己亲手搭建的框架——庭辩一旦开始，她只能在自己那个下午裁定的基底框架之上，应对对手的反驳、法官的追问。她不能在中场休息时偷偷掏出手机让AI替她“重新生成一版更优方案”。她被锁在了她自己的裁断里，必须在那个裁断的不可退回中，一路走到底。

这个设计并不试图制造“大量的”痛苦。一个下午的时间足够完成合格的准备工作，它并不比传统闭卷考试更消耗体力。但它精准地恢复了一样东西：在一切真正开始之前，那个最初的基底方案，必须是学习者自己、在没有AI代偿路径的条件下、亲手裁定的。这个裁断一旦做出，就成为后续全部

表现的不可退回的前提，学习者必须在它的约束下承受整个庭辩过程的全部冲击。这个设计，没有取消AI——学生在平时训练中仍可大量使用AI进行探索性学习；它只是在这个特定的、被制度标识为“关键承担窗口”的节点上，切断了退回路径，逼出了那个无法被代偿的意志动作。

这只是一个具例，不是唯一道路。不同的法学院、不同的人才培养定位、不同的课程类型，应当设计出各自形态的“不可退回承担窗口”。关键变量是同一个：是否在学习路径中，嵌入了至少在几个关键节点上不能被学习者单方退回的承担？如果是，教育仍在锻造决断主体。如果不是，无论产出指标多高，教育都在向培训漂移。而设计这些窗口的艺术，正在于它不是靠增加学生的痛苦来折磨人，而是通过自觉限制AI的在场，人为制造出学习者与学习对象（一个疑难案件）之间那种“没有任何中介缓冲”的直接磕碰。在那样的磕碰中，焦虑与清醒、挫败与洞见、以及最重要的——关于“我自己会如何选择”的不可逆认知，才可能一并发生。

六、可裁决性声明与结语

本文的核心论证，最终落座于一个严格的可被检验与反驳的命题之上。这个命题的内容是：在AI将法学教育核心训练环节中的不可退回承担系统性地撤销之后，学习者默会承担的发生——即那种通过亲手切断可能性而同步生成的、关于“我在此世选择了何种不可化解之价值紧张”的身体性认识——将无法在目前可观察的典型教学环境中有效发生。该能力的缺失，不会在学习阶段被充分暴露（因为它恰好被AI辅助下产出的高质量学业文本所遮蔽），而将在学习者此后进入高度不确定、不可逆且需要即时做出个人决断的法律实务场域时，才以“受过充分训练、却无法在压力下自主承担决断”的形态显现出来。

证伪条件如下：如果我们在未来广泛地、可重复地观察到，在系统性依赖AI辅助完成法学教育全部核心训练环节的法科毕业生群体中，其在上述高压、不可逆实务场域中的临场决断质量、价值承担的坚定程度、以及在犯下判决错误后主动承担而非将其归因于“当时AI分析有局限”的意志力，与从未接触过AI辅助的同龄对照组相比，并未出现统计学意义上的显著衰减——那么本文关于“逼视时刻所要求的不可退回性结构条件不可被撤销”的核心假设即被推翻。那将意味着，人类的默会承担的发生路径远较本文所设想的更具可塑性，“被退回的焦虑”或许确实可以经由其他尚未被本文识别的发生机制得到等价补偿。如果这一证伪成真，本文愿意成为它的一个注脚。在那一刻到来之前，本文坚守它所提出的这个令人不适的判断。

法学教育当前站在一个没有回头路的岔口上。一条路是继续拥抱AI带来的全部效率与舒适，同时将“不可退回的承担”从默认设置降格为选修的高级自律，然后安静地接受自己在未来被重新定义为一套高级法律决策培训系统的结局。另一条路是重新发明一套在AI的技术条件下，依然能够合法地、不可绕行地制造出与工具的对抗性遭遇的制度形式与教学设计——不是作为对旧时代的乡愁，而是作为它之所以是“教育”而非“培训”的那条唯一界线的再确认。这场发明目前尚未真正开始。本文的全部写作，只希望它开始得还不算太晚。

深化增补·全球近邻补切（编辑增补）

说明：以下各节为编辑在审读时增补，不改动作者正文与核心判断，只补上本文原稿未正面处理的最近邻理论，并把分离线写成可操作的判别。每一条都给出“若观察到什么，本文即错”的对照预测。

与卡普尔「生产性失败」的分离线：失败之后有没有那个「随后」

本文论证被 AI 代偿掉的不是艰难本身而是艰难内置的不可退回性。这一步很关键，但它使本文与学习科学里最贴近的那个传统正面相邻：卡普尔（Kapur, 2008）的生产性失败。他的实验结论是，让学习者先在无指导条件下失败地折腾，其长期迁移优于直接教授——「艰难有价值」在这一支里已经有实证。

分离线在于失败之后是否存在一个整合环节。生产性失败的效果依赖于「随后」：折腾之后必须有教师的整合性讲解，否则效果不成立，这在后续的对照研究中被反复确认。本文所说的逼视时刻恰恰是**没有随后**的——不可退回性的意思就是那一下之后没有人来收拾，承担在切断的同一瞬间生成。

对照预测：为一组学生保留完整的无指导折腾并配以充分的事后整合讲解（生产性失败的标准配方），为另一组保留少量真正不可退回的裁断。若两组在数年后面对无先例情境时的决断意志无差异，则本文的不可退回性条件多余，生产性失败足够；若有差异且差异只出现在「无人收拾」的那一类经验上，则本文成立。

把「合法草图的制度化」接到自动化偏见：一组现成的可测指标

本文正文已引用自动化偏见一线的研究（Skitka 等 1999；Cummings 2004），但只作为支撑提及。这里把它提到方法论位置上，因为它能让本文的核心命题变成可采集的量。自动化偏见研究给出的两类经典指标是**遗漏错误**（系统未提示时人也不发现）与**委任错误**（系统提示错误时人跟着错），并可测量核查行为的频次与深度。

把这套指标搬到法学教育现场，本文的「合法草图的制度化」就有了操作定义：在起草与检索任务中，测量学生对 AI 输出的独立核查率、发现植入错误的比率、以及在提交前修改自己原有判断的方向（向 AI 靠拢还是向自己的理由靠拢）。本文预测的不是这些指标随 AI 使用而恶化——那已经是自动化偏见的既有结论——而是**更进一步的一件事**：即使在完全离线、无 AI 可用的考场里，长期习惯于合法草图的学生仍然表现出决断的推迟与对不可撤回结论的回避。这一条与自动化偏见的分离线因此非常干净：偏见是在场效应，本文说的是离场之后仍在的效应。若离线条件下两代学生的决断行为无差异，则本文的装置闲置命题不成立。

本次增补所核验的文献

Kapur, M. (2008). Productive failure. *Cognition and Instruction*, 26(3), 379–424.

Loibl, K., Roll, I., & Rummel, N. (2017). Towards a theory of when and how problem solving followed by instruction supports learning. *Educational Psychology Review*, 29, 693–715.

Parasuraman, R., & Manzey, D. H. (2010). Complacency and bias in human use of automation. *Human Factors*, 52(3), 381–410.

参考文献（编辑按正文实际引用补列）

原稿正文以著者一年份的方式引用了若干文献而未附文献表。此处补列可确切核实者；另有 Bloch (2011)、Barry et al. (2000) 两处，因正文未给出足以定位的信息，留待作者补全，编辑不代为推断。

Bourdieu, P. (1980). *Le sens pratique*. Minuit.

Braverman, H. (1974). *Labor and Monopoly Capital*. Monthly Review Press.

Frank, J. (1930). *Law and the Modern Mind*. Brentano's.

Cummings, M. L. (2004). Automation bias in intelligent time critical decision support systems. AIAA Intelligent Systems Technical Conference.

Skitka, L. J., Mosier, K. L., & Burdick, M. (1999). Does automation bias decision-making? *International Journal of Human-Computer Studies*, 51(5), 991–1006.