

认知的禁食：AI时代课程功能的重释

课程的功能不再是提供信息，而是有纪律地不提供——并守住那段不被填补的时间

阳涌 著 · SDE 学员 · 课程理论与AI时代教育 · 2026年7月26日

摘要 当生成式AI可以在任何时刻、为任何人、以适配其当前水平的语言解释几乎任何知识时，课程作为“知识传递者”的旧身份便从底部被掏空。教育界的普遍回应是让课程转向“更高阶”的东西——批判性思维、创造力、元认知——但这个转向方案共享一个未经审查的前提：课程的本质仍然是一个传递装置，只不过传递的对象从知识换成了能力。本文拒绝从“课程应该是什么”出发去定义新本质。它转向追踪一个正在发生的事实：在AI已经可以即时回答一切的教室里，一些教师开始刻意地、有原则地“不回答”——不是因为没有办法，而是判断此刻给出答案会跳过学生认知发展所必需的那段“自己在混乱中摸索”的间隙。本文把这种正在发生的教学位移命名为“认知禁食”，并通过生成机制追踪、概念不可还原性校验、历史谱系的重构、操作边界的划定以及对教育不平等风险的正视，展开对这一概念的论证。认知禁食不是AI应激的权宜之计，也不是课程被发现的“真正本质”——它是课程在信息丰裕化条件下一轮尚未完成的生成中被逼出来的新显态，其制度性壳体（评价的社会授权、身体在场的羞耻与信任、制度性不可退出）构成了它在结构中不可替代的那部分。最后，本文给出证伪条件，并正面回应认知禁食可能加剧教育不平等的反驳。

关键词：认知禁食；课程功能；生成式AI；信息丰裕化；productive failure；教育不平等

一、课程为何物？一个被问错的问题

有一件事正在发生，而我们用来描述它的词汇已经不够用了。

这件事是：一个学生可以在午休时用手机问清楚光合作用的卡尔文循环，在写论文的间隙让AI帮他梳理论证结构，在复习备考时获得一份为他量身定制的、比任何教材都更清晰的知识梳理。如果课程的核心功能只是把已知的东西讲清楚，那它就被一个装在口袋里的设备全面替代了。

面对这种冲击，教育界的反应是迅速而焦虑的。各国的课程改革方案、学校的教学创新计划、教育科技公司的宣传文案，都在用类似的措辞描绘出路：既然AI能教知识，课程就该教那些“更高阶”的东西——批判性思维、创造力、协作能力、元认知。这个转向的共识之广，以至于它听起来已经不像是一个判断，而像是一条不证自明的公理。

但它不是。把“批判性思维”写进课程目标，不等于课程真的知道怎么培养批判性思维。更根本的问题在于：这些方案共享一个未经审查的前提——课程仍然是一种传递装置，只不过以前传递的是知识，现在要传递的是“能力”或“素养”。这不过是把货物从一艘正在下沉的船搬到另一艘看起来更坚固的船上，却没有问一句：也许这个海运业本身已经过时了。

这就是为什么关于“AI时代课程怎么办”的讨论，尽管声势浩大，却始终在一层薄薄的技术更新话语上打转。更新内容、更新方法、更新工具——这些工作都重要，但没有人追问一个更前的问题：当知识不再需要被“传递”时，课程在做什么？

这个问题本身需要被重新审视。因为它悄悄预设了“课程当然存在，只是需要重新定义”。但也许课程之所以还在，不是因为它的某个更深本质终于被我们发现，而是因为它在被迫做一件它从前不需要刻意去做的事——这件事在传递功能被抽空之后，反而被放大为它必须直面的事。

本文试图描述的，正是这件事。不通过定义课程的新本质——那种做法假设答案是一个可以被一次性命中、然后固定下来的“本质”——而是追踪一个正在发生的位移：在AI可以即时回答一切的教学现场，一些教师开始刻意地、有原则地“不回答”。他们不是没有答案，而是判断此刻给出答案会跳过学生认知发展所必需的那段间隙。这个“故意不回答”的行为，许多教师已经在日常教学中默默实践，但它还没有被命名，更没有在理论上被严肃地追问：它是否可能是课程在AI时代被迫承担的新功能？它能否被任何已有理论框架中的概念完全覆盖？它在历史上是否有过类似的制度化形态？它在落地时面临哪些边界和风险？

这个追问需要做六件事。第一，解剖课程的旧身份——知识传递者——建立在什么样的认知生态之上，以及AI如何把这个生态的根基抽空。第二，追踪一个当代教学冲突的现场案例，让“认知禁食”从理论推演落到真实的生成现场——不是“应该这样做”，而是“已经在这样发生”。第三，执行严格的概念锻造工序，从课程与AI的矛盾深处逼出那个不可还原的新命名，并正面碰撞教育心理学内部最危险的最近邻——Kapur的productive failure——完成不可还原性校验。第四，回溯这个新功能的历史先例，但拒绝把不同时代的发生条件不同的实践焊接成一条“本质谱系”。第五，划定认知禁食与最容易被混淆的三种实践——知识库、训练营、算法推荐——之间的操作边界。第六，论证课程的制度性壳体——评价授权、身体在场的羞耻与信任、制度性不可退出——是这个新功能唯一能在结构中落地的底

座。最后，在结语中正视认知禁食被滥用的历史事实和加剧不平等的风险，给出证伪条件，让整个论证保持向反驳开放。

这篇文章不为“课程怎么办”提供确定性方案。确定性方案是另一种做法的终点。我们正站在一轮尚未完成的生成之中，目前最诚实的事，不是假装已经看到了终点，而是把这场生成本身说清楚——它的张力在哪里，它的不可替代性在哪里，它可能失败的条件又在哪儿。

二、认知生态的位移：AI如何瓦解了课程旧身份的运转条件

要理解课程正在经历什么，必须先理解它曾经赖以成立的条件。不只是“信息不对称”这个社会条件，而是在信息不对称庇护下，学习者的认知装置是如何被日常学习节奏反复激活的。“知识传递”这个说法太轻易了，它把一整个认知生态的复杂运作压缩成一个从A到B的输送动作。我们必须把这个压缩包解开。

在信息不对称时代，学习者的认知运作实际上依赖一个隐性的“输入瓶颈”。这个瓶颈不是谁设计的，而是技术条件的自然结果：能接触的知识就那么多，能回答你问题的人就那么几个，能查到的资料就那么些。这个瓶颈看似是限制，实则承担了一个关键的认知功能：它强迫学习者在信息不足条件下，调动已有的知识结构和判断力去填补空白、做出推测、形成暂时的结论。换句话说，信息的稀缺逼着人去主动处理那仅有的一点信息。当你手里只有一本教材、一个老师和三本参考书时，你自然而然地会把它们翻来覆去地读，在脑子里反复组合它们给出的线索，试图拼出一个你自己认为合理的整体图景。这个“自己拼”的过程，就是那种被称为“认知自制力”的习性在日常学习中被反复激活的训练场。

这里必须小心措辞。我使用“认知自制力”一词，不是指一种先天的、独立于情境的心理能力，而是指一种在特定认知生态中被反复调用而养成的习性：在面对一个陌生问题时，不立刻寻求外部成品答案，而是先调动自己已有的东西去试探性地组织一个框架，并在试探过程中容忍“不确定”和“暂时站不住”的不适。它不是某种可以被直接教给的技能，而是一种只能在“亲自处理信息”的过程中慢慢长出来的心理倾向。

AI没有消灭知识——恰恰相反，它让知识变得前所未有的丰富、易得、可定制。但这不是问题的关键。关键是：AI消灭了那个曾经让认知自制力在日常学习中被反复激活的“输入瓶颈”。当一个学生可以在遇到任何认知卡点时即时获得一个结构清晰、解释充分、甚至适配他当前水平的答案时，他的认知装置就不再被逼着去“自己填补空白”。他可以跳过那个最难受、也最关键的处理阶段——在信息不全的情况下试探性地组织理解——而直接获得成品。输入瓶颈的消失，让信息处理从“主动拼”滑向了“被动接收”。不是AI让他变懒了，是他的认知习性在缺少必要摩擦的环境中，失去了被激活的条件。

这就可以解释一个令人困惑的现象：为什么AI明明让学习变得更方便了，许多教师却报告说学生的深层理解能力反而在下降？不是因为AI教错了什么——AI教的内容往往比教师更准确、更全面。而是因为AI的“方便”恰好填平了认知习性生长所必需的那种让人不舒服的间隙——从遇到一个问题到获得一个答案之间，那段不知道该怎么、必须自己去试一试的时间。那段间隙，就是认知自制力被反复激活的生长窗口。它需要的不只是认知上的挑战，还有一种心理上的“忍不住想跑但被某种力量按住了”的体验。在过去，信息输入瓶颈自然地制造了这种“被按住”的处境——你找不到现成答案，只能自己硬扛。在AI时代，这个“被按住”需要被刻意制造。课程的角色位移，正发生在这里。

这也就解释了为什么“课程升级教批判性思维”这条路走不通。批判性思维不能被当作另一个“东西”来传递——把批判性思维的知识拆解出来、编成模块、做成练习题，这种做法恰恰是把批判性思维当成了旧课程中的“知识”来对待。它的根本问题在于：它仍然假设课程的核心操作是传递——只不过传递的东西从内容换成了技能。但AI也能传递批判性思维的框架、步骤、技巧和范例，而且往往比教师传递得更清晰、更系统。如果课程只是“传递批判性思维”，那它同样会被AI替代。

认知自制力不能被传递——它只能在被逼着亲自处理信息的过程中长出来。它不是一种可以被言传的知识，而是一种在反复经历“信息不足→推测→碰壁→修正→再推测”的循环中慢慢养成的心理习性。这种习性一旦养成，就能迁移到任何新领域：拥有这种习性的人，在面对一个陌生问题时，不会立刻寻求外部答案，而是会先调动自己已有的东西去试探性地组织一个框架，然后在试探过程中发现哪里站不住、哪里需要补充。AI可以给他任何答案，但AI给不了他这种不急于要答案的耐心。这种耐心不是性格上的温顺，而是认知上的自律——知道真正的理解必须经过一段没有现成答案的混乱期，并且愿意待在那段混乱里。

课程旧身份的瓦解，从这个角度看，就不仅仅是“信息不对称消失导致传递功能丧失”。更深的那一层是：AI通过消除输入瓶颈，同时消除了认知自制力在日常学习中被激活的那个生态环境。课程旧身份的核心运作，不是它传递了知识，而是它在传递知识的过程中，不自觉地附带制造了输入瓶颈——那些学生等不到答案、必须自己先想一想的时刻。AI把输入瓶颈摧毁之后，课程就暴露出它真正在发生位移的部分：不是传递，而是在信息供给过剩的环境中，刻意制造和守护那些正在消失的认知间隙。这就是新功能从旧身份瓦解的矛盾里被逼出来的方式。

但这里必须做一个关键的区分。“刻意制造认知间隙”如果不加限定，就可能滑向一种危险的主张：教师可以随心所欲地隐瞒信息，学生越困惑越好。这不是本文要论证的。新功能真正内核不是“制造困难”——困难有太多可以滥用的空间——而是在信息供给过剩的

环境中，保护认知自制力的生长窗口不被即时满足的习惯所关闭。它的目标不是让学生痛苦，而是让学生重新获得那种“在没有现成答案的情况下，用自己的判断力撑住一段不确定”的习性。这个目标与“制造困难”有本质的区别。前者是手段，后者可能变成目的本身。本文将在后续操作边界的划定中严格锁定这个区别。

三、一个生成现场：研讨课上的“不回答”

上一节完成了对课程旧身份瓦解的诊断。但诊断不等于发生。生成机制要求追踪的不是“应该发生什么”，而是“正在发生什么”——在真实的、有摩擦的教学现场里，课程的功能位移是如何被具体的冲突和应对逼出来的。

这一节提供一个当代研讨课的案例。案例来自一所研究型大学的社会理论研讨课。它不是我亲历的田野调查——本文属理论建构，不基于实证研究——而是对一类越来越频繁出现的教学情境的典型化重构。这类情境的共同特征可以这样描述：一门以深度阅读和课堂讨论为核心的研讨课，在AI工具普及之后，教师既成事实论生的讨论发言越来越流畅、越来越“对”、但越来越没有真正的判断力。他们能用AI在课前把文本“读懂”——总结出核心论点、梳理好论证结构、给出支持与反对意见——然后在课堂上复述这些整理好的东西。但他们无法在追问中用自己的话重新解释文本的核心概念，无法在两种看起来都对的解读之间做出有理由的选择，也无法识别AI生成的论证中隐含的前提错误。

我以这个情境为基础，重构一个具体片段的展开过程。

教师布置了阅读文本——一篇论证密度很高的理论文章。在之前的学期，教师通常会在课前给出导读问题，帮学生进入文本。这学期，他做了一个调整：不再给出导读问题，而是要求学生在课前提交一个问题——必须是他们自己在阅读中“真正想不通的、没法用一句话绕过去的问题”，而不是那种可以轻松回答的“作者说了什么”。他在第一堂课上明确告诉学生：“你可以用任何工具来帮助理解文本。但你在课前提交的问题，必须是你自己在和文本死磕之后仍然卡住的东西。如果AI能替你想出这个问题，那它就不是你的问题。”

这句话本身，就是一次认知禁食的发生现场。教师没有禁止AI——AI就在那里，谁都可以用——但他通过“提交一个自己的问题”这个教学设计，在信息无限供给的环境中切出了一块AI无法代劳的间隙：AI可以总结文本，可以生成可能的问题，但它无法辨认哪个问题是你自己在理解中真正撞上的那堵墙。那个“撞上”的体验，是你的，不是AI的。教师通过这个要求，把学生从“获取AI生成的成品”逼回了“自己先磕一遍”——不是因为他认为AI不好，而是因为磕的过程本身，就是认知自制力被激活的生长窗口。

课堂上，一个学生提交的问题是：“作者在第二节说X，在第四节说Y——我怎么觉得X和Y是矛盾的？但我知道作者肯定不是傻子，所以可能是我没读懂。”这是典型的“认知间隙”——学生已经隐约感觉到了文本内部的张力，但他的第一反应是怀疑自己，而不是怀疑文本。如果他在此刻用AI去问“作者第二节和第四节是不是矛盾了”，AI会给他一个条理清晰的分析，指出矛盾只是表面的，深层的逻辑是这样的——然后他就永久失去了自己亲手拆解这个张力的机会。

教师没有回答他的问题。不是因为他也不知道，而是因为他判断这个学生此刻正站在一个认知发展的关键节点上：他已经隐约看到了矛盾，但他还需要自己走进这个矛盾，把两边的论证重新组织一遍，看看它们到底在哪个假设上分叉了，为什么在作者那里这两个假设可以并存。教师只说了一句：“你的直觉是对的。你现在不要查AI，也别问我。你回去把第二节和第四节的论证前提各写一句话，然后告诉我，这两句话能不打架的条件是什么。下节课你来回答。”

这就是一次有原则的“不回答”。它的结构可以分解如下：（1）学生在认知上已经准备好处理这个矛盾——他感觉到了矛盾，只是还能力厘清它。（2）教师基于对学生的了解和他的认知发展阶段做出了判断——如果现在给他现成答案，他会绕过那个厘清的过程。

（3）教师的中断是暂时的、可逆的——他明确给出了下节课的回应时限和具体操作路径（写两句话、找出不打脸的条件）。（4）学生不能退出这个“不回答”——他要在下节课面对全班和教师，给出他自己的回答。

这个案例展示的是：认知禁食不是理论家的构想，是教师在AI时代被逼出来的应对。教师不是从某篇论文里学到“要做认知禁食”——他是在具体的教学困境中，发现旧的教学设计（给导读问题、在课堂上讲解）已经在AI环境下失效了：学生看似准备好了，但那个“准备”是AI代劳的。为了逼出学生自己的认知处理，他必须改变教学设计——不是教更多，而是有原则地不给。这个“不给”不是懒，不是不负责，而是在新的认知生态中对教学责任的重新理解：保护学生认知自制力的生长窗口，比给他一个正确答案更重要。

这个案例也将为后续的概念锻造提供经验锚点。在下一节，我将从这个案例中抽象出认知禁食的操作结构，然后通过概念碰撞检验它是否是一个不可还原的新辨别维度。

四、认知禁食：一个被矛盾逼出来的新命名

前面两节完成了两件事：第一，从认知生态位移的角度论证了课程旧身份的瓦解；第二，通过一个具体的教学案例展示了新功能如何在AI时代的课堂现场被逼出来。现在要做的事，是给这个正在发生的事一个精确的、不可还原的命名。不是造一个新词，而是从矛盾的结算里逼出一个任何已有概念都装不下的新辨别维度。

4.1 从案例中抽象操作结构

从上一节的研讨课案例中，可以抽象出“有原则地中断信息供给”的四个结构要素。这四个要素共同构成了这个概念的操作内核。

第一个要素是认知势能的识别。教师之所以选择“不回答”，不是因为懒，而是因为识别到学生已经具备了处理这个矛盾的基础条件——他已经自己读出了X和Y之间的矛盾，只是还能力厘清它。这个识别是认知禁食的起点：没有它，“不给”就退化为随意的藏匿。它要求教师对特定学生在特定时刻的认知状态有足够精准判断——这个判断是专业性的核心。

第二个要素是暂停供给的决策。教师拒绝回答学生的问题，不是因为没答案，而是判断此刻给出答案会跳过学生认知发展的关键节点——他自己厘清矛盾的那个过程。这个决策的逻辑不是“信息不好”，而是“时机不对”。在学生已经形成了自己的追问、但还没尝试自己回答之前，外部成品答案会把那个追问堵回去。暂停供给，是为了让追问在学生的脑子里继续发酵。

第三个要素是可逆的时限与路径。教师的“不回答”不是永久性的知识封锁。他给出了明确的时限——“下节课你来回答”——和明确的操作路径——“把论证前提各写一句话，找出不打脸的条件”。这个可逆性是认知禁食区别于知识管控的分水岭：门是关的，但学生知道门后面有东西，也知道怎么去敲那扇门、什么时候会开。

第四个要素是不可退出的关系结构。学生不能绕过这个“不回答”。他要面对的不是一个可以点击“跳过”的界面，而是一个下节课就要见面的、他尊敬的老师，以及一个他每周都要面对的同辈群体。这个结构性的“不能退出”，不是权力的粗暴行使，而是认知禁食得以施加于那些会自动回避不适的头脑上的前提条件。没有这个结构，暂停供给就只是建议——而建议在即时满足的文化中几乎没有约束力。

这四个要素的拼合，构成了一个已有的教学设计概念无法完整覆盖的结构。下一小节将通过与最危险的最近邻——Kapur的productive failure——的正面对质，来检验这个结构的不可还原性。

4.2 正面对质：认知禁食与productive failure的分离点

在所有的理论邻居中，最危险的——也是最不该回避的——是Manu Kapur的productive failure研究。Kapur及其合作者通过一系列对照实验证明：在学习新概念的初期，如果让学生先在没有直接指导的条件下尝试解决他们尚不能完全解决的问题、生成多样性的方案并在尝试中遭遇失败，之后再接受直接教学——这种“先失败后教学”的序列，比“先教学后练习”的传统序列更能提升深层理解和知识迁移能力。这个操作逻辑与认知禁食在表面上有极强的同构：两者都主张不在第一时间给出指导，两者都强调让学生自己先“磕”一阵，两者都诉诸认知摩擦对深层学习的价值。

这个同构性是真实的，但正是在这里，概念分离才成为可能。如果认知禁食只是productive failure在教学实践中的另一种说法，那就是可还原的——应该撤回“认知禁食”这个概念，改为“productive failure的制度性实施条件研究”。

分离点不在于“不给指导”这个操作层，而在于这个操作被嵌入的结构关系。

让我以四问法逐一推进。

第一，Kapur谈论的是什么条件？Kapur的productive failure研究，关心的是任务序列的设计条件：在什么类型的任务上、以什么方式分组、在什么样的时间分配下，“先尝试失败再做直接教学”比“先教学后练习”更有效。他的控制变量是任务序列和指导时机，他的结果变量是概念理解和迁移。他偶尔会提及班级文化和社会规范的影响，但这些不在他的理论硬核之中。他的理论硬核是认知—任务的。

第二，认知禁食多追问的条件是什么？认知禁食讨论的，恰好是productive failure的边界条件本身：在什么制度性条件下，那种要求学生先自己磕一阵的教学设计，能真正施加在那些会自动回避不适的学生身上？如果学生可以随时打开AI拿到答案——或者更根本地说，如果学生可以在不适出现的第一秒就退出这个任务——那么再精妙的任务设计也没用。认知禁食指向的，是让“不能退出”成为可能的那套关系、制度与伦理条件。它的理论硬核是制度—关系的。

第三，两者的变量不在同一层。用一个类比：productive failure研究的是“什么样的运动训练计划能提升耐力”，认知禁食研究的是“在什么制度—关系条件下，一个习惯了坐电梯的人愿意每天爬楼梯”。前者关心训练内容的设计，后者关心训练内容施加于人的前提条件。两者不矛盾，但不可互相还原：你可以有世界最好的训练计划，但如果受训者可以随时退出，训练就不会发生。在AI环境中，可以随时退出是默认条件——这正是productive failure研究在诞生时没有面对的认知生态。Kapur的实验预设了学生身处一个不能轻易退出的课堂——但今天这个预设本身正在瓦解。认知禁食追问的，正是这个预设瓦解之后，课程如何重新建立“不可退出”的条件。

第四，两者的伦理关切点不同。productive failure的伦理基础是“这样做学习效果更好”——它属于教学效果的范畴。认知禁食在此基础上多了一层伦理可问责性：教师决定此刻不给答案，不是因为“不给答案的学习效果更好”这一认知规律（虽然它确实援引这个规律），而是因为这个决定本身是一个可以被追问、被质疑、被修正的专业判断。他可以也必须在事后面对学生、同事、家长的追问：“你

当时为什么选择不给答案？”如果他的回答是“因为按照productive failure的研究，这样效果更好”——这是一个教学技术判断。如果他的回答是“因为我就在那个时刻判断，你已经有能力自己迈出这一步，只是需要被推一把——而如果不推，你可能永远不会自己试”——这是一个嵌在人际关系里的认知—伦理判断。后者不能被前者覆盖。

因此，认知禁食与productive failure的分离线可以这样画：productive failure研究的是任务层面的认知序列设计；认知禁食研究的是使这些设计在“即时退出是默认选项”的认知生态中仍然能够落在学习者身上的制度—关系条件。认知禁食的操作当然包含productive failure的认知逻辑——在这个意义上两者是互补的，认知禁食可以引用productive failure来为“暂时不给”提供认知依据——但认知禁食多出了一个productive failure不追问的维度：那个“不给”的执行者，是一个嵌在制度壳体中的、对特定学习者的认知发展承担责任的具体的、他做出的每一次“不给”，都不仅要回答“这样有效吗”，还要回答“我对这个学生的判断对吗？我对得起他对我的信任吗？”

这组分离线，使得认知禁食在概念上不可还原为productive failure的教学论版本。它不是对Kapur的否定，而是对Kapur框架在AI时代认知生态中制度性边界条件的追问和补全。

在下一节，我将通过与卢曼、伯恩斯坦和比约克的碰撞，进一步确认认知禁食作为一个复合体概念的不可还原性。同时，我将给出这一概念的精确操作定义。

4.3 精确定义与剩余近邻碰撞

基于案例的抽象和与productive failure的对质，现在可以给出认知禁食的操作定义，然后通过与三个剩余近邻的碰撞来收紧它的边界。

认知禁食：一种在信息可以无限即时获取的技术条件下，由可问责的制度主体（教师/课程）基于对特定学习者认知状态的判断，主动、可逆、有路径地暂时中断外部信息的即时供给，从而逼出那种在没有现成答案的条件下用自己的判断力去试探性组织理解的习性——的制度化实践。

这一定义的每一个限定词都有重量。“可问责的制度主体”——不是任何人的随意行为，而是承担教育责任的教师，其决定可以被专业共同体和学生审查。“基于对认知状态的判断”——不是普遍规则，而是针对特定学习者的专业判断。“可逆、有路径”——不是永久封锁知识，而是指明开启的条件和时限。“逼出试探性组织理解的习性”——目标不是让学生痛苦，而是让那个在即时满足环境中日益失去激活条件的认知习性重新获得生长空间。“制度化实践”——不是个别教师的个人风格，而是可以被检验、被争论、被修正的课程操作。

现在用三场简短的对质，把认知禁食与三个容易混淆的已有概念之间的边界廓清。

与卢曼（社会系统免疫论）的分离。卢曼用免疫隐喻分析社会系统如何通过排斥环境侵扰来维持自身边界的封闭性和自创生能力。在他的框架里，教育系统排除某些知识——比如不把占星术纳入课程——不是某个人做的决定，而是系统自我再生产时自动伴生的功能：系统“免疫”于那些会瓦解其自身再生产逻辑的信息。这个排除是无人称的、功能性的，系统不为具体的排除行为担责。认知禁食与卢曼免疫论的根本分离点，在于“可问责性的制度化”。认知禁食的每一次“中断供给”，不是系统的自动反应，而是一个具体的教师基于对具体学生的判断做出的决定——而且这个决定必须能说出理由、接受质疑和修正。一个教师决定这节课不给出答案，先让学生自己撞半小时——这个决定可以被同行评议：你是不是高估了学生的先备知识？你是不是在回避某个自己讲不清楚的知识点？你是不是用“制造困难”来掩盖教学准备不足？这些问题是对一个可被问责的人提出的。卢曼装不下这个伦理可问责性维度。

与伯恩斯坦（教育知识编码）的分离。伯恩斯坦论证了课程对知识的选择、排序和组织从来不是中性的，而是社会权力再生产的隐秘机制。某些知识被纳入核心课程、某些被边缘化——背后是阶层再生产的逻辑。伯恩斯坦框架如果应用于认知禁食，会把它批判为一种更精致的守门行为：你以为你在保护认知发展，实际上你在用一套更进步的词汇，把某种特定的知识等级制度包装成认知必然性。分离点在排除的时间性与目的。伯恩斯坦分析的守门行为，是社会结构性的、期待门永远关着的排除——工人阶级的语码被挡在正式课程之外，不是“等学生准备好了再放进来”，而是被永久宣判为在学术场域中没有位置。认知禁食的排除逻辑，是发展性的、阶段性的、以将来放进来为前提的。在基础阶段不讲量子力学，不是因为它低级，而是因为判断这个阶段引入它会阻碍而非促进深层理解——等到经典物理框架稳立之后，量子力学不仅会被引入，而且会被引入为对经典物理的一个根本挑战，让学生在两种范式的碰撞中重新审视自己曾深信的东西。守门人的排除期待门永远关着；认知禁食的排除恰恰在等待门被撞开的那一刻。伯恩斯坦的理论无法解释这种“为了将来放进来而此刻关上门”的操作逻辑，因为在他的分析框架里，课程排除的伦理前提永远是“永久”。

与比约克（理想困难）的分离。认知心理学中比约克夫妇的“理想困难”概念证明：那些让学习在短期内更困难的条件——间隔练习、交错练习、变异练习——反而提升长期记忆保持和迁移能力。认知禁食的核心操作——在信息供给过剩的环境中刻意制造认知摩擦——与理想困难共享一个地基。但认知禁食多出一个理想困难不处理的维度：那个“困难”得以被施加于学习者的制度—关系结构。理想困难是任务层面的认知特征——间隔、交错、变异——它们可以被编进算法。事实上，自适应的智能辅导系统已经在精准实施理想困难：AI可

以根据学习者的表现自动调整间隔、交错不同题型的练习、增加变异性。AI在这个意义上可能比人类教师做得更精准。但什么不能被编进算法？是那种让学生不能点击“跳过”的制度—关系结构。学生在AI辅导系统上遇到一道让他极度不适的难题时，他可以关掉窗口、切换到娱乐内容、或者让AI直接给出答案。在课堂上，当一个学生被你尊敬的老师当众追问“你刚才说的那个前提，你真的相信吗”——你不能关掉窗口，不能换个语气重说，不能假装网卡了。你被钉在那种不适里，而正是这种无法逃脱的不适，逼你去做AI界面无法逼你做的事：在羞耻和自尊并存的状态下，重新审视自己刚刚说过的判断。理想困难告诉你“要设计什么样的任务特征”；认知禁食告诉你“仅仅设计任务特征还不够——还必须有一种拒绝退出的结构，才能让理想困难真正施加在那些习惯于回避不适的头脑上。”这个拒绝退出的结构，就是课程的制度性壳体。这个概念将在第六节被充分展开。

三场对质之后，可以确认：认知禁食的不可还原部分，不在于它的“排除”（卢曼有），不在于它的“制造困难”（比约克有），也不在于它“基于认知发展阶段的判断”（与productive failure共享），而在于它把这三种成分嵌进了一个可问责的制度关系里——在这个关系里，排除是暂时的、以将来发生认知冲突为前提的，困难是嵌入在不可退出的社会—伦理结构里的，并且每一个中断供给的决定都要为自己的判断接受检验。这个多重限定结构的复合，使得认知禁食不可被任何已有单一概念无损重述。

五、历史先例的重构——不是谱系，而是家族相似

上一节完成了认知禁食的概念锻造。这一节要处理一个概念锻造之后常见的陷阱：把这个新概念投射回历史，把不同时代、不同条件、不同生成路径的教育实践焊成一条“跨历史的谱系”，好像它们都是同一个永恒本质在不同时代的显影。

我不会做这个。我更倾向于另一种做法：诚实地审视历史上那些与认知禁食在操作层面有“家族相似”的教育实践，然后追问它们各自是在什么条件下被生成出来的。追问的目的不是证明“认知禁食古已有之”，而是借历史中的先例来理解：在信息环境发生结构性变化时，教育制度可能会被逼出什么样的反应。这些历史先例是参照，不是祖宗。

5.1 研讨班：一次对信息丰裕化的制度性回应

有一个历史先例与认知禁食的发生条件最为接近。十八世纪末，印刷术的成熟使得书籍成本大幅下降、图书贸易繁荣，大学生手里已经有了教授在课堂上念的那本书——甚至它的更新版本。随之而来的是对大学讲授课的激烈质疑：如果学生已经能看到这些文本，讲授课还有什么存在必要？

普鲁士各大学——尤其是洪堡改革后的柏林大学——给出的回应，不是提高讲课质量。那只是把旧传递升级。他们发明了研讨班。在研讨班里，教授不是把知识加工好递出去，而是反其道而行之：他把学生扔进原始材料里——法律文本、历史档案、古典语言残篇——逼他们自己去分解、排序、组织。然后在他们自以为理清的时候，用追回去拆他们搭起来的东西。研讨班所制造的，正是一种被精心设计的“暂时不给现成解释”。它不给学生那个教授已经知道的路标，逼他们从不给中自己生出解释；然后不给这个解释太平落地，再把它搅乱，逼出更稳的解释。

研讨班的生成条件可以这样描述：在信息获取不再困难之后，教育制度被逼着把“信息处理重新变难”——不是通过降低信息质量（研讨班用的是一流原始文献），而是通过切断“跳过我自己的处理就能拿到成品”的路径。在讲授课里，知识的分解和排序是教授完成的，学生接收的是加工过的成品；在研讨班里，学生必须自己去分解和排序，教授不帮他做这一步，反而在他做完之后用追问把它弄乱，逼他再做一次。这个操作结构与认知禁食共享同样的内核：在外部信息可以轻易获得成品的条件下，刻意在通往“成品理解”的路径上设置一段只能由学习者亲自穿越的间隙。

5.2 三种家族相似的实践——但生成条件各不相同

研讨班之外，还可以列举三种在操作上与认知禁食有家族相似的历史实践。但必须预先声明：它们各自的生成条件与研讨班完全不同，不能把它们装入“认知禁食谱系”这个大筐。

其一是中世纪修道院的圣言诵读。修士面对一段经文，不是去查阅注疏，而是被要求反复诵读、默想、在心中来回咀嚼同一段文本，直到它自己的层次自己展开。这个操作确实包含“不给现成解释、逼你自己先磕”的成分。但它的生成条件是信息极度稀缺——修道院不是有更好的解经书而故意不给，而是解经书还没被写出来、或者极难获得。换言之，圣言诵读中的“不给”，不是丰裕中的刻意中断，而是稀缺中的不得不给得慢。它生成于一个与认知禁食完全不同的信息生态。

其二是文艺复兴时期人文主义学校对经典文本的背诵。学生被要求用母语和拉丁语交替口述同一段论证，逼他们在两种语言的转换中重新组织逻辑结构。这个操作也包含“不让你直接拿到成品翻译”的成分。但它的生成条件——虽然处于印刷术开始传播的时代——根植于修辞教育的古典传统，而非对信息丰裕的制度性回应。它的目标不是刻意中断信息供给，而是通过语言转换训练来培养雄辩能力。

其三是东方传统的经典诵读。在被允许接触注疏和讨论之前，学童被要求把核心文本背诵到成为一种身体记忆。这个操作包含“先不

给你解释、先背”的成分。但它的生成条件同样与信息丰裕无关——它根植于对经典的崇敬和记忆在口传文化中的中心地位。

这四种实践——研讨班、圣言诵读、人文学校背诵、东方经典诵读——在操作表面都包含“暂时不给现成解释”的成分，这让它们在事后看起来像是同一个本质的家族。但它们的生成条件截然不同。研讨班是被印刷术时代的信息丰裕逼出来的制度创新；圣言诵读是信息稀缺时代的不得已的慢；人文学校背诵和东方经典诵读则分别根植于修辞传统和经典崇敬。把它们归入同一条“认知禁食的跨历史谱系”，就是用新发明的筛子去retrospectively归类历史——那是既成事实论的典型操作。本文不做这种归类。本文尊重它们各自的发生条件。

5.3 历史先例给今天的启示

那么，这些历史先例对今天有什么启示？不是“认知禁食古已有之，所以它是对的”，而是两条更具体的认知。

第一，当信息环境发生结构性变化时，教育制度确实会被逼出某种“暂时不给”的回应。研讨班是迄今为止最接近的制度化先例——它的生成条件（印刷术导致信息丰裕化）与认知禁食的生成条件（AI导致信息丰裕化）在结构上同型。研讨班的命运因此具有参照价值：它成功了——研讨班从洪堡时代一直延续到今天，成为研究型大学的核心教学形式。但它的成功不是因为“它发现了教育的本质”，而是因为它所在的制度壳体——研究型大学的社会授权、教授与学生的长期指导关系、学位制度的不可退出性——提供了它运作的结构条件。认知禁食如果要制度化，同样需要这样的壳体。

第二，认知禁食的落地，必须严格限定在“信息丰裕化”这一生成条件内。它的操作前提不是“信息稀缺，只能不给”，而是“信息过剩，反而必须刻意不给”。这个区分至关重要：如果认知禁食被解释为“在任何情况下都应该先不给答案”，那它就会在信息稀缺的教室里——先备知识薄弱、家庭支持匮乏、学生最需要明确的指导和充足的信息时——退化为以发展之名行剥夺之实。这个风险，本文将在结语中正面回应。

六、在三个邻居的夹缝中：认知禁食的操作边界

一个新概念能从矛盾里被逼出来、有历史的参照，还不够。它必须在操作上能与最容易被混淆的三种实践——知识库、训练营、算法推荐——划清边界，否则它就会在落地时滑回旧范式。这一节要逐一确立这些边界。

6.1 与知识库的边界：基于认知判断的暂停，而非不提供

知识库——无论是传统的百科全书还是动态的AI问答系统——的伦理是完备性：给用户一切相关的、可用的、能帮助理解的信息。一个AI助手的默认状态是“你问什么我答什么”。它不以说“不”为美德。

认知禁食恰恰以说“不”为核心操作——但它说的“不”必须严格限定在特定条件下执行，否则就与“不提供”没有区别。不提供是知识库的缺席——是资源的缺失、是入口的封闭、是学生想要却被挡住。认知禁食是暂停供给——是在具备完全信息能力的条件下，基于一个可被检验的教学判断做出的临时性决定。这二者有截然不同的认知后果。

不提供让学生意识不到自己缺少了什么东西——他不知道自己被蒙在鼓里，因此也就不会产生认知饥饿感。暂停供给的前提，恰恰是学生已经隐约意识到这里有更多的东西，但目前还不能一次性接收——他正在撞上自己的理解边界，正是这个“撞上”的体验让他产生了想要跨越边界的动力。不提供是封闭式的——墙永远在那里；暂停供给是开放式的——门是关的，但学生知道门后面有东西，而且他知道这扇门会在某个阶段被打开。

这就是认知禁食第一个操作边界：它不主张课程变成一个处处说不的封闭系统，而主张在特定认知节点上做出有意识的、可解禁的延迟。它不是知识库的低配版，它是知识库在经过教学判断之后的选择性降速——降速的目的不是让学生永远不知道某些东西，而是让他们在知道之前，先长出能消化这些东西的认知结构。

6.2 与训练营的边界：保留不可优化的弯路，而非反效率

编程集训班、考前冲刺班、领导力工作坊——这些训练营与认知禁食共享一个表面特征：都有明确的目标、结构化的步骤、以及某种程度的“不给现成答案”。但训练营的“不给”是战术性的——不给答案是为了让你更快掌握解题模板；当你掌握了模板，那个“不给”就自动撤销了。训练营的效率哲学是：所有认知摩擦最终都应该被优化掉。

认知禁食保留的摩擦恰恰是不可优化的——不是因为做不到，而是因为优化掉它，就等于优化掉了认知自制力生长的唯一土壤。一节好的数学课让学生用牛顿定律从头推导开普勒的行星运动定律——这个过程三百年前就有人做完了，学生这辈子可能也不会再推一次。但在这条弯路上，他必须把力的分解、向心力、微积分、守恒律这几个概念亲自拉起来协调运转。他会在某个步骤撞上：所有环节单独看都成立，拼在一起却推不下去。他必须回头找自己的假设哪里不对。当他终于推通的那个时刻，他获得的不是正确公式，而是一种“我亲手

把一个复杂的东西捏合拢了”的秩序感。这种秩序感，就是认知自制力的一次激活。

训练营会直接给他最佳推导路径。AI会给他每一步的详细解析。它们都节省了时间，但剥夺了那个唯一能激活认知自制力的间隙。认知禁食与训练营的根本分界线就在这里：训练营追求效率最大化，认知禁食保护那些不能被效率侵蚀的认知发展过程。它不反对效率——在教学的许多环节，效率是正当且必要的——但它拒绝让效率成为唯一仲裁者。

6.3 与算法推荐的边界：逆着即时满足的惯性，而非反对个性化

算法推荐——短视频、搜索建议、AI个性化推送——在诱惑力上有课程无法匹敌的结构优势：它永远给你此刻最可能愿意点开的东西。它的伦理是让你舒服地留在信息流里。课程天生就是反算法的——它推荐的下一个知识，是你此刻最需要消化的，不是你此刻最想要的那种舒服的知道。算法在你刚有朦胧疑问的时候就排好相关答案送过来；课程在你刚有朦胧疑问的时候，很可能把这个疑问悬在那里，不给答案，让你自己再去撞、再去想。

认知禁食不是“反对个性化”。恰恰相反——好的认知禁食必须建立在高度个性化的认知诊断之上。因为暂停供给什么时候做、暂停多久、暂停到什么程度，取决于对这个特定学习者在此时此刻的认知状态的精准判断。暂停早了，学生会因为完全摸不着边而放弃；暂停晚了，“已经知道”就是答案了，没有禁食的必要。认知禁食与算法推荐的真正分界，不在个性化程度的高低，而在个性化服务的伦理方向：算法推荐是用个性化帮你绕开不适，认知禁食是用个性化把你精准地安置在你此刻最需要的那一种不适中。

这也就回答了为什么认知禁食不能是一个批量化的课程设置——它不能是“本课程所有章节均不提前给出结论”这种一刀切的规则。一刀切的“不给”恰恰最容易被滥用，因为它省去了教师对每个学生此刻认知状态的判断——而这个判断，正是认知禁食的伦理核心。没有判断的不给，不是禁食，是怠惰。这条分界线，将在下一节关于制度性壳体的讨论中获得制度层面的支撑：认知禁食之所以不是教师的个人任性，是因为它必须被嵌入在一套以发展性判断和专业可问责性为核心的操作边界之内。越过这些边界，“不给”就退化为隐瞒与支配。

七、壳体的重铸：课程如何在制度生活中撑起认知禁食

前面三节解决了认知禁食是什么、从哪里来、与邻近概念如何划界这三个问题。但还悬着一个最关键的问题：在AI也能拒绝给出答案、也能制造认知摩擦的未来——也许很快——课程的认知禁食还有什么不可替代？这个问题不能靠拉高嗓门再把旧论点喊一遍，必须去审视课程在制度性存在中所嵌入的那些结构性事实——那些AI在原则上无法复制的东西。

第一件事，是评价的社会授权。课程不只是教学系统，它同时是社会授权的价值判断装置。一门课给你分数、评语、证书——这些不仅是学习反馈，它们向社会发出信号：这个人经过了一段被制度化监控的认知训练，被一个受外部审核的机构认可为合格。这个认可之所以有分量，不是因为任课教师个人更聪明、更有判断力，而是因为他嵌入在一整套制度里。他的评价标准要经过同行审议，他的课程要接受质量评估，他给出的学分在劳动力市场上有法定效力和交换价值。

这个制度授权构成认知禁食全部“拒绝退出”结构的地基。学生在认知禁食操作中不能随意退出，不只是因为“暂时性的认知势差”，更是因为退出有真实的社会代价——不通过、不毕业、不被认可。这个代价不是认知层面的，是社会契约层面的。AI可以给出比任何教师都精准的认知诊断和困难设置，但它没有签发社会承认的作者权。这不是技术问题——再强的AI也不能单方面改变学历认证法律。认知禁食的真正力度，不全在“禁食”设计得有多高明，而在于它背后有一套不能轻易退出的制度安排。而这一整套安排，AI在可预见的未来都无法获得，因为它需要的不是更强的智能，而是社会契约的重写。

第二件事，是身体在场的羞耻与信任。在AI对话中暴露自己的认知漏洞也会产生挫败感，但那是私人的、没有目光重量的。课堂上被一个你尊重的人当众追问时，暴露的不只是你不会某道题，而是你的整个判断习惯、你回避问题的倾向、你在压力下为自己找借口的那些细微动作被真实地看见了。这种“被看见”产生两种AI无法生成的东西。

其一是羞耻——逼你做AI纠错时你不需要做的那种深度自我检视，因为你被看低的不是一道题，而是你这个人此刻的判断质地。这里必须小心。羞耻在教育中的使用充满危险。心理学家对羞耻与内疚做了关键区分：内疚指向行为——“我做错了这个”——而羞耻指向自我——“我是错的”。前者倾向于激发修补行为，后者倾向于引发逃避或攻击。课堂上那种“被追问的不适”，如果滑向羞耻——学生觉得自己整个人被否定了——就会产生相反效果：他会逃避这门课、憎恨这位老师、或者用敷衍的方式消解被追问的压力。但如果这种不适被限定在“你的论证这里站不住”而非“你这个人不行”，它就更接近一种“认知性问责压力”——一种与对特定推理行为的不适紧紧绑在一起的、指向修补和重检的压力。认知禁食需要的是后者，不是前者。教师让学生“下节课你来回答”时传递的信息是“我相信你有能力回答，所以我暂时不替你做”——这是一份被信任包裹着的压力。如果学生接收到的信号是“你不行，所以你活该被晾着”，那不是认知禁食，那是精神虐待。

其二是信任。在与一位老师经过长期磨合之后，你开始相信他给你的批评不是否定你，而是他知道你能承受这个批评并且会在上面长出新东西。这种信任只能在时间中、在双向的、有风险的相遇中被赢得——不是被编程，不是被模拟。只有在这种被真实他者见证的关系

里，那种认知性问责压力才会转化为重检判断的推力，信任才会转化为在困难中不退出的定力。缺少这一层关系质地，认知禁食就退化为算法化的理想困难——可以被跳过，可以被忽略，因为对面没有一个你在乎的人。而这两样东西都要求一件事：身体在场。不是“实时视频”意义上的影像在场，而是“同一个物理空间”意义上的、我躲不掉你的看见、你躲不掉我的反应的在场。这种在场在教育中意味着什么，至今没有AI能够模拟。不是因为技术不够，而是因为“在场”本身不是信息传递的表征——它首先是一种共同暴露的处境：你和我同时暴露在彼此的目光与判断里，无所遁形。这种暴露对互动的质地——尤其是它施加在个体身上的那种“必须回应、无法重来”的重量——无法被还原为算法对情绪线索的识别与反馈，无论那种识别有多么灵敏。

第三件事，是制度性的不可退出。这是课程嵌入的最深的不对称性。当你和AI对辩时，你握有离开的最终权力。但一门课——一旦你选了、一旦你进入、一旦你开始接受评价——你就无法说退就退。不是技术上不能，是制度上：退了就必须承担不通过、学分缺损、毕业延迟、社会信号失落这些真实后果。

这个不对称性，过去常被批判为教育的非人化。但当信息环境把“退出不适”优化到只剩一个按钮时，正是这个被长期诟病的特征，获得了全新的教育价值：它是唯一还能有效地把学习者强制留在认知困难面前的制度性力量。它不是暴力式强制，而是框架性保护——像护栏，你知道你可能会在这段路上摔，但护栏保证你不会摔出这条路。在一个人工智能、算法推荐、即时数字娱乐共同把“退出不适”优化到只剩一个按钮的世界里，学校制度里剩余的那些“不能退”——不能随便换掉这个太难的任务，不能跳过这段必须硬撑的泥泞，不能因为今天不想面对就不面对——变成了一种独特的、不可替代的认知保护机制。

这三种结构性事实——评价的社会授权、身体在场的羞耻与信任、制度性的不可退出——共同构成课程作为认知禁食的制度性壳体。壳体之内，AI可以是强大的辅助工具，可以在禁食的间隙里提供精准的个性化挑战和反馈。但壳体本身，AI没有——不是暂时没有，是原则上无法通过更强的技术获得。因为这些不是信息处理功能，是制度生活。认知禁食之所以是课程的功能而不是AI的功能，不是因为课程更聪明，而是因为课程有能力拒绝学习者退出。这个能力，来自制度，来自关系，来自对学习者的未来的社会契约责任。只要课程还承担着这份责任，发生在它内部的每一次认知禁食，就带有一个AI模拟不出的维度：它被一个真实的人在一个被授权的场域里逼出来的、不可退出的、因此可能真正改变认知习性的体验。

八、禁食的风险——一个结语，而非结论

如果一篇论文主张某种教育操作，它就必须在同一页纸上正视这个操作可能被滥用、被扭曲、被用来伤害它声称要帮助的人的所有方式。认知禁食的历史前身——从研讨班的精英排他性到古典背诵的服从训练——足以敲响警钟。这一节要做的，不是“回应对手反驳”的辩论套路，而是把这种危险内化为理论自身必须容纳的边界条件。

8.1 当“不给”成为剥夺

认知禁食面临的最根本的伦理风险，在于它的理论结构——基于认知发展阶段的判断来暂时中断信息供给——在操作中极易滑向对弱勢学习者的系统性剥夺。

一个中产家庭的学生和一个低收入家庭的学生，面对同一个“暂时不给答案”的教学操作，实际处境可能截然不同。前者在晚餐饭桌上可以向高学历的父母请教，可以在周末的补习班获得补充讲解，可以用家里的私人AI助理获取答案的替代版本。后者的“不给”就是真的不给——没有补充资源，没有替代路径，只有一扇关上的门。对前者来说，那扇关上的门可能通向一间更大的书房；对后者来说，它可能只是一堵墙。

这不是理论瑕疵，而是一个结构性风险：认知禁食的操作前提——“学生已经在某个认知发展节点上，暂停供给会逼他跳一跳”——在班级情境中需要进行针对每个学习者当下状态的诊断，而不仅是作为一揽子设计的前提。当一个教师用同样的教学设计面对二十五个知识基础、语言能力、家庭支持各不相同的学生时，“暂时不给”很可能不是“制造适当的认知摩擦”，而是在不具备必要前提的个体那里叠加上制度性的剥夺。他还没有足够的内部资源去承受这个“不给”——他不是“不想跳”，他是脚下没有能站的地方。

这不是说认知禁食不能做。而是说，认知禁食的合法操作，必须与补偿性条件配对出场：在被中断供给的学生面前，必须同时有足够的知识支架——不是给他答案，而是给他足够的地板，让他有地方站——以及明确的信任关系，让他知道自己不是被抛弃，而是被挑战。没有这些条件，“禁食”就不是禁食，是断食。而断食对饿着的人比对饱着的人更危险。

8.2 可逆性：不是程序，是伦理

在第三节的定义中，“可逆”是认知禁食的限定词之一。这个限定词不能只停留在概念声明层面——它必须被操作化。可逆不是一个时间表上的“下周三打开”，而是学生能够感知到并相信的东西。

让一个中断供给变成可逆的，不是教师单方面说了“下节课开放讨论”，而是学生在被中断的期间内，始终有途径获得最低限度的支

持——能问“我这个方向对吗？”而不会被训斥“让你自己想你还问我”；能在感到完全迷路时得到一个最小的路标，而不是被沉默地放任在挫败中。在这种条件下被中断供给的学生，才会把“不给”理解为教师对他的判断——“你在那个时刻能自己接住一定程度的混乱”——而不是被惩罚。可逆性保的不是时间表，是认知禁食的伦理正当性：教师不可以为了“锻炼学生”而把他们弃置在无法承受的挫败中。

8.3 证伪条件

最后，这篇论文必须给出自己的证伪条件。这不是学术礼仪，而是对自己所论证的东西能给出的最诚实的回答。

本论证把课程的不可替代性定位在它的制度性壳体——评价授权、身体在场的羞耻与信任、制度性不可退出。这三者共同构成了认知禁食没有退化成“可选困难”的结构原因。

因此，认知禁食作为课程专属功能的论证，在以下条件下被证伪：如果未来的智能学习系统能够同时满足三项——在制度层面获得社会授权签发学历与资格；在关系层面让学习者对其产生真实的、不可随意换掉的羞耻与信任；并在结构上建立学习者不能随意退出的制度性不对称——那么本文的核心论证就被证伪，认知禁食就不再是课程的专属功能，而成为任何满足这些制度条件的智能系统都可能执行的功能。

是否真的会有这么一天，我不知道。但把这个条件写在这里，至少让这篇论文不容忍自己变成一个自我免疫一切反驳的封闭命题。

认知禁食不是课程的救赎。它是课程在信息丰裕时代被逼出来的一种正在发生的位移——张力重重，边界模糊，充满被滥用的风险。把它写成论文，不是为它唱赞歌，是试图把这场位移说清楚。说清楚之后，它能不能在制度中站稳，能在多大范围内不被扭曲地落地，能不能不变成新一轮教育不平等的精致借口——这些不是论文能回答的。它们取决于一线教师在每一次“不给”时对特定学生的判断、取决于制度在每一个环节对这份判断的支持或压制、取决于我们是否愿意诚实地追问：在今天这个什么都可以即时获取的世界里，教育中那些让人不舒服的间隙，是真的在保护什么，还是仅仅在延续什么。

如果这篇以“认知禁食”开场的论述，到最后能让人感到的仍是那个“究竟在保护什么还是延续什么”的张力，那它就没有在结论中背叛它的生成机制起点。

注释

[1] 本文属理论建构，未采用田野调查或实证研究。第三节出现的教学案例为对一类在AI时代越来越频繁出现的教学情境的典型化重构，并非特定真实个案的报告。案例中的人物、对话与课程设置均为论证之便而设。

[2] 本文对卢曼免疫论的借用仅限于其“系统通过排斥维持边界”的功能类比，不采纳其理论中的“系统无意识”与“功能必要性”预设。卢曼社会系统理论的原始论述见Luhmann (1984)。

[3] 本文仅借用伯恩斯坦“分类”与“架构”的概念工具来分析知识排除的形式机制。伯恩斯坦关于教育知识编码与社会再生产的原始论述见Bernstein (1971)。

[4] 比约克关于理想困难的经典论述见Bjork (1994)。对理想困难边界条件的进一步讨论及元认知干扰效应的研究，见Bjork & Bjork (2011); Soderstrom & Bjork (2015)。

[5] Kapur关于productive failure的系统阐述见Kapur (2008, 2012, 2016)。本文对productive failure与认知禁食的分离开启于如下事实：Kapur的框架以任务序列设计为其理论硬核，不处理执行“先不给后教”的教学设计所需的制度—关系条件。认知禁食因此不是对productive failure的替代或否定，而是对其在AI时代认知生态中的制度性前提条件的追问——这一追问不在Kapur文献的问题域内。

[6] 第五节所涉历史实践（研讨班、圣言诵读、人文学校背诵、东方经典诵读）均为基于通行教育史叙述的参照性例示，本文不将其作为经验证据使用，更不将其建构为与认知禁食在发生条件上统一的“跨历史谱系”。第五节的核心论证在于：这些实践各自的发生条件不同，共享的只是操作表面的家族相似；仅研讨班是在“信息丰裕化”条件下被生成出来的制度性回应，因此与认知禁食的发生条件最为接近。关于研讨班的历史，参见McClelland (1980)关于德国大学与国家关系的论述；关于圣言诵读，参见Leclercq (1982)关于中世纪修道院学习文化的经典研究。

[7] 关于羞耻与内疚在心理学中的区分及其对学习行为的差异化影响，参见Tangney & Dearing (2002)的综述。本文使用的“认知性问责压力”概念并不完全对应Tangney区分中的任何一类，而是试图在教育关系中描述一种指向特定推理行为而非整体自我的、以修补和重检为导向的不适体验。

参考文献

- Bernstein, B. (1971). *Class, Codes and Control, Volume 1: Theoretical Studies towards a Sociology of Language*. London: Routledge & Kegan Paul.
- Bjork, R. A. (1994). Memory and metamemory considerations in the training of human beings. In J. Metcalfe & A. P. Shimamura (Eds.), *Metacognition: Knowing about knowing* (pp. 185–205). Cambridge, MA: MIT Press.
- Bjork, E. L., & Bjork, R. A. (2011). Making things hard on yourself, but in a good way: Creating desirable difficulties to enhance learning. In M. A. Gernsbacher et al. (Eds.), *Psychology and the real world* (pp. 56–64). New York: Worth.
- Kapur, M. (2008). Productive failure. *Cognition and Instruction*, 26(3), 379–424.
- Kapur, M. (2012). Productive failure in learning the concept of variance. *Instructional Science*, 40(4), 651–672.
- Kapur, M. (2016). Examining productive failure, productive success, unproductive failure, and unproductive success in learning. *Educational Psychologist*, 51(2), 289–299.
- Leclercq, J. (1982). *The Love of Learning and the Desire for God: A Study of Monastic Culture*. New York: Fordham University Press.
- Luhmann, N. (1984). *Soziale Systeme: Grundriß einer allgemeinen Theorie*. Frankfurt: Suhrkamp.
- McClelland, C. E. (1980). *State, Society and University in Germany, 1700–1914*. Cambridge: Cambridge University Press.
- Soderstrom, N. C., & Bjork, R. A. (2015). Learning versus performance: An integrative review. *Perspectives on Psychological Science*, 10(2), 176–199.
- Tangney, J. P., & Dearing, R. L. (2002). *Shame and Guilt*. New York: Guilford Press.

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：与合意困难、生成效应的硬边界

正文与 productive failure 划了分离点，但学习科学里另有两个更近的邻居未被处理：比约克的合意困难（提取练习、间隔、交错等在当时降低表现却提升长期保持的操作）与生成效应（自己产生答案优于被给予答案）。分野在被节制之物：合意困难与生成效应节制的是学习条件的舒适度，其目的与判据都在最终掌握上——若长期保持没有提升，该困难就不合意；认知禁食节制的是信息供给本身，且其价值不完全由掌握结果兑现——那段不被填补的时间可能不产出任何可测的掌握增益，而它守护的是学习者与未解问题独处的经验。可判别面：若禁食的全部效应都能被长期保持指标吸收，则它只是合意困难的一个变体；若在长期保持无差异的条件下仍出现自主提问率与问题重定义能力的差异，本文获得独立地位。

增补二：禁食的操作化与安全边界

禁食若无操作定义与边界，极易退化为教师不作为的托词。操作定义：在学习者明确提出信息请求之后，教师有意延迟提供的时长，及该时长内是否提供了任何替代性支架（提示、缩小范围、同伴讨论）——真正的禁食是不提供替代支架的延迟。安全边界有三条硬线：不适用于安全相关知识；不适用于学习者已出现明显焦虑或退出迹象时；不适用于评价即将发生的时段（否则禁食的代价由学生的成绩承担而非由课程承担）。三条边界必须写在操作定义之前，否则本文会成为一种把责任转移给学生的教学法。

增补三：家族相似而非谱系的方法论说明

正文以家族相似而非谱系处理研讨班等历史先例，这是一个正确且需要被明确辩护的方法选择。辩护如下：谱系主张要求证明制度间的实际传承关系，而研讨班、书院、师徒制的生成条件各不相同，强行串成一条线会把不同的因果机制混为一谈；家族相似只主张这些实践在某一项操作上同形——都以纪律性的不提供来制造一段不被填补的时间。这一较弱的主张足以支撑本文的论证，且它明确排除了一种诱人的说法——认知禁食是某个古老传统的复兴。它不是复兴，它是在新条件下被重新逼出来的。

增补四：与作者已发表工作的分野

作者已发表《当教育删除了知识被生成的那个瞬间》与《合法性的牢笼：学校认证如何将不可见的学习逐出课堂》，两篇处理的是那个瞬间被删除、那部分学习被逐出，其分析位置在损失一侧；本文的位置在供给一侧——课程如何有纪律地不供给，从而把那段时间守住。分野可一句话说清：那两篇问被拿走了什么，本文问该扣住什么不给。两者构成同一判断的诊断与操作两面，宜互相标注。

编辑增补所依据的核验文献

Robert A. Bjork & Elizabeth L. Bjork, “A New Theory of Disuse and an Old Theory of Stimulus Fluctuation,” in *From Learning Processes to Cognitive Processes*, Erlbaum, 1992.

Manu Kapur, “Productive Failure,” *Cognition and Instruction*, 26(3), 2008.

Norman J. Slamecka & Peter Graf, “The Generation Effect: Delineation of a Phenomenon,” *Journal of Experimental Psychology: Human Learning and Memory*, 4(6), 1978.

Jerome S. Bruner, *Toward a Theory of Instruction*, Harvard University Press, 1966.

Lev S. Vygotsky, *Mind in Society*, Harvard University Press, 1978.