

活壳：数字时代成就悖论的发生学分析

任何因高度适应外部而稳定下来的生存结构，在成就自身的同时都潜在地抽空了那个使它诞生的活的转化力

作者 张琼 · SDE 学员 · 约 2.5 万字 · 发表于2026年7月31日 · 深化增补版

摘要 数字平台赋予个体空前的外部能力，却催生了一种普遍的生命体验：人越来越“能干”，却越来越难感受到“做成了一件事”的踏实。既有批判理论将这一悖论归咎于外部力量的侵蚀，但它们共享一个未经审查的预设：那个被侵蚀的主体内核依然完好，只需移除外部障碍便可恢复生机。本文论证，这一预设本身正是问题的根源。通过撤销“主体性必须靠亲自处理外部世界来生长且不可被结构性替代”这一更深层的先验，本文从当代数字生存的矛盾中结算出一个此前未被命名的新结构：活壳——任何因高度适应外部而稳定下来的生存结构，其在成就自身的同时，都潜在地抽空了那个使它诞生的活的转化力；当抽空越过临界线，“成就”与“坏死”便不再是两件事，而是同一过程的两张面孔。本文从这一结构的发生学机制入手，追溯了活壳如何在每一次“这次可以绕过去”的分叉选择中逐步结晶，并在与延伸心灵论、斯蒂格勒无产阶级化理论等最相近理论的判决性对质中划清分离线。本文还以一个完整案例展示了活壳从支架滑向假体的全过程，以自身为对象进行方法论自反，公开其证伪条件，并指出出路不在“回归真我”，而在承认坏死的前提下，去探索一种从未存在过的生成方式。

关键词 活壳；悖论性成就；生成性转化力；假体宿命；锤炼绕过

一、引言：一个尚未被命名的现象

我们正活在人类能力史上一个奇特的巅峰。一个从未受过专业写作训练的年轻人，借助大型语言模型，可以在一小时内产出结构完整、引经据典的万字长文。一个零工经济的参与者，仅凭一部智能手机，便能调度跨越半个城市的服务链条。如果我们将“能做成一件事”的标准定义为一个动作被启动、一个过程被完成、一个结果被产出，那么数字时代的个体正处于历史能力的顶峰。

然而，与这种外部能力急剧膨胀同步发生的，是一种隐晦却普遍的内心塌陷。一个越来越常被提及却难以精确描述的感受是：人们觉得自己越来越“能干”，却越来越难感受到那种真正“做成了一件事”的踏实感。你每天都在处理海量信息、完成大量对话、产出大量内容，但在关闭屏幕的那一刻，一种空转感涌上来——你好像做了很多，却又好像什么都没做。不是你不够努力，而是你的努力似乎沉淀不下任何东西。

这种感受不是抑郁，不是倦怠，不是传统意义上的“意义危机”。它是一种更根本的、发生在行动与存在之间的断裂：在一个行动效率被无限放大的时代，行动者本人却感到自身行动能力的萎缩。滑动的次数越多，进入深度关系的能力越弱；产出的文本越高效，独立组织论证的能力越萎缩；被算法推着走的“劲头”越充沛，内在驱动自己找方向的那种真正的劲头就越枯竭。我把这种成就与空转同源发生的现象，称为成就-坏死的一体两面性。它不是两个互斥的现象碰巧同时出现，而是同一个结构同时产出两种产物。

面对这一状况，现有的主流解释集中于三种进路，但它们各自留下了无法覆盖的盲区。

第一种进路将问题归结为“选择过载”与“注意力碎片化”。它抓住了“量”的冲击，但其逻辑终点是：只要个体学会更好的注意力管理，问题就能缓解。然而，许多高度自律、擅长管理注意力的高效能人士，仍然深陷那种空转感。问题似乎不在注意力的分配，而在于行动本身的质地已被悄然置换。

第二种进路来自监控资本主义批判传统。肖莎娜·佐博芙揭示了数字资本如何将人类经验转化为可交易的行为数据，借此预测和塑造人的行动，侵蚀其自主性。这一“被夺走”的叙事极具力量，但它隐含一个乐观前提：被夺走的东西，可以被夺回来。它暗示着通过抵抗、立法、技术伦理改造来“收复失地”的可能。然而，如果问题不只是“被夺走”，而是“被替代”——如果那件被拿走的东西在原位上被安上了一个好用的假肢，而身体在长期依赖假肢后，自己长不出真肢了——那么夺回就成了一句空话。

第三种进路援引早期批判理论，尤其是赫伯特·马尔库塞在《单向度的人》中提出的诊断：发达工业社会通过技术理性和消费主义，提供了一种“压抑性去升华”，用即时的满足卸载了本可指向反抗与超越的势能。它准确刻画了张力被提前释放的过程，但它预设了势能的源头仍在，只是被消费主义疏导了。可如果势能的源头本身正在枯竭——如果一个人不是因为欲望被满足而平静，而是因为他身体里产生欲望的那块肌肉已经萎缩了呢？

这三条批判路径——选择过载、监控掠夺、压抑性满足——共同分享了一个未被言明的前提：它们都假定那个能选择、能反抗、能被压抑的主体内核依然完好。问题被描述为主体遇到了坏的外部环境，只要移除或修正这些外部因素，内核就能恢复生机。

本文要论证的恰恰相反：在长期的、高效的外部功能延伸之下，那个内核本身已经发生了结构性的退化。不仅如此，本文将进一步指出，这一退化揭示了一个此前未被命名、却在无数生命领域中反复发生的结构。我将其命名为“活壳”。

活壳的核心判断是：任何因高度适应外部而稳定下来的生存结构，其在成就自身的同时，都潜在地抽空了那个使它诞生的活的转化力；当这一抽空越过临界线，“成就”与“坏死”便不再是两件事，而是同一过程的两张面孔。一个人依靠外部模板产出了漂亮的论文，这是一种成就；但他独立从混沌中理出秩序的能力，在这种成就中被静悄悄地坏死了。一个平台依靠算法推荐维持了用户的高活跃度，这是一种成就；但用户内部那个生成方向的原初驱力，在这种成就中被逐步替换为外部信号的机械反射。成就越大，壳越硬；壳越硬，里面越空。

在正式展开之前，我需要做一笔方法论自反。引言部分对既有批判进路的分类——选择过载、监控掠夺、压抑性满足——在形式上是发现学的写作姿势：先给既有解释分类，再指出它们共享的盲区，然后提出自己的新框架。这一结构是任何一篇学术论文在引言阶段都无法完全绕开的必要妥协。我清楚自己在这里做的是发现学的动作，但我不会在后文中继续使用这种“分类-填补”的论证方式。后文将不把这些批判进路当作需要被否弃的错了解释，而是把它们当作同一片土壤中长出的不同显露态

——活壳要追溯的，不是“它们哪里错了”，而是“这片土壤本身是什么，以及它是如何把成就与坏死作为同一过程的两张面孔一并结算出来的”。

对本文反复指涉的“亲自处理”“亲自承受阻力”“从混沌中生成秩序”等概念，也需预先做一笔澄清：本文对“亲自处理”的反复强调，不是对其价值体系的怀旧式重申，而是因为——在当前的人类生物条件下，某些核心转化力的生成恰好需要一个“亲自承受某种特定类型的阻力”的锤炼过程。如果未来人类的生物基底发生了改变（例如通过神经接口技术使得生成性自我效能可以通过非锤炼途径建立），那么活壳的坏死诊断可能需要被修正。本文的论域严格限制在“揭示发生学机制”，而非“论证亲自比被承接更有价值”。价值的判断是读者自己的事。

二、活壳：一个结构的显影

我不打算在本章开头给出“活壳”的定义。那个定义应该是在分析完成之后、作为一种被逼出来的结晶而落在纸上。我们现在要做的事，是撤掉一个无人质疑的底层先验，然后进入一个具体的矛盾，看那个结构如何从矛盾里一点一点地长出来。

2.1 撤到底：那个连批判者自己都还站着先验

要凿出一个真正不可还原的新结构，不能靠“又多想到一个角度”，而必须从撤销一个无人质疑的底层先验开始。让我们对成就-坏死的一体两面性，执行一次逐层深入的追问。

第一问：这个领域里最痛的矛盾是什么？是人越来越能干，却越来越虚空。外部能力的延伸与内部根基的收缩同步发生。

第二问：旧的解释路径为什么反复失效？因为它们都将问题归咎于外部——信息过载、资本掠夺、压抑性满足——却从未追问：那个“被过载”“被掠夺”“被压抑”的主体，其内部是否已经发生了某种更根本的退化？它们开出的药方——管理注意力、抵抗资本、解放欲望——都预设了那个主体内核依然完好，只需要被唤醒或保护。但这种唤醒为什么反复失败？

第三问：这些失效背后，大家默认为真、从来没有质疑过的那个先验假设是什么？是：只要外部条件改善了，内部能力就能恢复。这是一个深植于现代批判思想骨髓里的信念。它假定内部能力的根基是自足的、不可被外部条件结构性改变的，外部条件只能“压迫”它或“解放”它，但不能“掏空”它。人有一个不变的“本真内核”，历史只是在这个内核外面加上了不同的枷锁或装饰。

这个先验还不够深。对它再追一刀：这个“不可被掏空的本质内核”预设本身，又预设了什么？

它预设了：主体性只能通过“亲自处理外部世界”这一种方式来生长和维持。无论你是在抵抗压迫、在管理注意力、在解放欲望，你都被假定为必须亲自面对世界、亲自在混沌中摸索、亲自在失败中锤炼。这套“亲自处理信念”是整个现代主体哲学的基石。从笛卡尔的“我思”到康德的“自律”，从杜威的“做中学”到批判理论的“觉醒”，都在不同层面重复着这个前提：主体性的唯一生成路径，是主体亲自去应对、去克服、去转化。

第四问：如果撤销这个先验，会打开什么？

会打开一个此前被遮蔽的维度：主体性也可以被外部结构“承接”走，并在这种承接中发生质变。不是被压迫，不是被掠夺，而是被温柔地、高效地、不知不觉地卸下了。一个学生不再亲自面对空白文档的焦虑，因为AI替他完成了从混沌到秩序的组织工作。一个用户不再亲自在无聊中慢慢升起一个属于自己的冲动，因为算法推荐替他填满了每一寸注意力空白。在这个过程中，那个“亲自处理”的功能没有被压迫——它干脆没有被征召。而一个长期不被征召的功能，不会只是“沉睡”——它会坏死。

连这个“坏死”的判断本身，又预设了什么？预设了坏死是不好的。这个预设隐藏着一个更深的价值判断：亲自生成的、从内部驱动的、笨拙但真实的东西，优于外部高效替代的、顺滑但空洞的东西。但凭什么？为什么“亲自”就一定比“被承接”更好？

这个问题逼出了一个最深的底层先验——人类文明迄今为止的整个价值秩序，都建立在“亲自处理信念”的基石上。我们赞美创造、赞美努力、赞美在挫败中咬牙前行，是因为我们默认了：人只有通过亲自去磕、去碰、去做，才能成为人。一旦承认了主体性可以被外部承接、并且在这种承接中“活得很好”——甚至活得比亲自去磕的人更顺滑、更快乐、更少痛苦——整个价值秩序就会动摇。

正是在这里，我站住了。我不再往下追，因为再追就会陷入“亲自处理信念”究竟值不值得捍卫的伦理之争，而那已经超出了本文的发生学关切。本文的任务不是论证“亲自”比“被承接”更好，而是揭示一个此前未被命名的结构：任何外部承接性结构，在成就一个主体的同时，都潜在地取代了那个使主体得以生成它自身的活的转化力；当这种取代越过某条临界线后，“成就”与“坏死”就不再是两件事，而是同一过程的两张面孔。

我把这个结构叫作“活壳”。它不是我先定义好、再拿去找例子的概念——它是从上述四重追问的终点上，被逼出来的命名。接下来的任务，不是继续定义它，而是回到一个具体的生命现场，看它如何在一次次选择中被结算出来。

2.2 逐次结算：一个微观发生学的追溯

活壳不是静态的“壳”，也不是一个先在的诊断框架。它是一个在无数次微小的分叉选择中，被一次次结算进个体的行动模式里的动态结构。本节要做的，不是给出一个“诱绕-断修-伪稳”的三步模型——那种模型是分析结束后的概括，而不是分析开始前的框架——而是进入一个真实的微观序列，在每一个关键的分叉点上，追问同一个问题：这一次选择的条件是什么？这些条件中哪些是被前一次选择的结果重新塑造的？下一次选择的空间又是如何被这一次逐渐收窄的？

以下追溯的主人公，是一个我称之为“林晚”的研究生。她的经历将在第六章被展开为一个完整的发生学案例。此处我仅借用她的故事中那些具有典型性的分叉时刻，来展示活壳的结算机制——不是在讲“她怎么了”，而是在讲“每一次选择是如何成为前一次选择的惯性产物的”。读者可以在阅读时自行检验：这些分叉节点，是否已经在自己的经验中反复出现。

第一次绕过：一个选择的分叉点

一个研究生在面对一篇需要独立完成的论文时，第一次尝试使用AI辅助写作。起初，她只是让它润色几个句子——这是支架式的使用：她先自己写出了论证，只是在语言层面寻求辅助。她在核对修改时仍在思考“为什么这样改更好”——锤炼依然在发生，只是被部分地辅助了。此处还没有活壳。

活壳的第一个分叉点出现在她面对一个困难的论证段落时。她感到一阵强烈的思维阻力——那种阻力是大脑真实承受认知负荷时的生理信号。在旧的学习模式里，她必须坐在这片阻力里，反复尝试、反复失败，直到从混乱中理出线索。这个过程痛苦、缓慢、充满自我怀疑，但正是这种痛苦在逼着她的前额叶皮层搭建新的神经连接。

但此刻，她有一个更轻松选项：直接把问题丢给AI，让它生成一个论证框架。她犹豫了几分钟。这个犹豫本身就是一个分叉点——两个方向的引力在同时拉她。一边是“我应该自己试试”的残余信念，另一边是“这太痛苦了，而且AI做得可能更好”的解脱诱惑。

什么条件在这次分叉中占了上风？至少三种条件在共同作用。第一，截止日期在逼近——时间压力放大了“快速完成”的权重，压低了“慢速锤炼”的吸引力。第二，上次她自己硬写的论文得了低分，而同学用AI辅助拿了好成绩——外部评价的历史反馈，已经在她体内积累了一种“绕过更划算”的隐性知识。第三，那个阻力区本身带来的生理性不适——前额叶在真实思考时的能量消耗是巨大的，大脑天然倾向于节能——让她在面对“硬扛”选项时产生了某种近乎身体的排斥。三种条件合在一起，把“犹豫”推向了“绕过”。

她把问题输了进去。AI在三秒内生成了一个结构清晰的框架。她看了一眼，觉得“差不多”，在此基础上修改了两处措辞，补充了一段自己的细读，提交了。分数不错，老师没发现异样。第一次绕过成功。

第五到第十五次：归因模糊的累积

下一次面对思维阻力时，那个犹豫缩短了。不是因为她做了有意识的决定——而是因为上一次的经验沉淀成了一条隐蔽的神经通路：“上次绕过之后，结果不差，而且少受了很多罪。”这条通路不是在信念层面被“记住”的，而是在神经层面被强化的——每一次绕过都伴随着多巴胺的微量释放（任务完成、外部肯定），使得“绕过”这个动作本身被正面强化。

到了第七八次，一个新的条件加入进来：同伴效应。她发现室友也在用同样的方式写论文，没有人觉得这有什么不对。社会参照系的引力进一步消解了“我应该自己试试”的残余信念——当周围人都在绕过时，硬扛不仅痛苦，还会让你显得笨拙、低效。

到了第十次，另一个更深的变化在自我叙事层面悄悄完成了。她开始在心底把自己描述为“那种不太会写论文的人”——而不承认的是，这个“不太会”并非天生的能力缺陷，而是她自己在过去半年里，每一次都选择了绕过锤炼之后，结算出来的一种功能退化。但她不知道自己的选择正在结算出这个结果，她以为这就是“我不行”。于是，“我是那种需要工具帮助的人”这个自我认定，在每一次绕过中被反复确认，最终内化为一种稳定的自我叙事。而这个自我叙事，又反过来为下一次绕过提供了最强大的合法性——“我本来就不会，所以用工具是合理的。”

临界时刻：空腔的第一次暴露

某一天，她面临一个没有网络、没有AI辅助的即时发言场景。她必须在十几分钟内，用自己的脑子组织出一个论证。她发现，那个曾经在本科阶段反复训练过的“从混沌中理出结构”的能力，此刻像一块太久没被征召的肌肉——她能模模糊糊地感觉到它应该在哪儿，但大脑的指令传达不过去。她开口说了几句，语无伦次，然后慌乱地放弃了。

那一刻，空腔第一次被她自己触摸到。

这个空腔不是一次“绕过失败”的结果——它是无数次“成功绕过”的结果。每一次绕过都是一次成功：更高效、更高分、更少挫败。但每一次成功的绕过，都在同一个节点上完成了一次微小的切除——那个节点就是“亲自承受认知阻力并从中理出

秩序”的核心锤炼环节。十五次成功，就是十五次切除。切到最后，那个被切除的器官不再是一个“休眠的功能”，而是一个崩解的结构——它不是因为太久没用而变生疏，而是因为它的神经通路已经被一条更高效的替代回路给彻底覆盖了。

这就是活壳的结算机制：不是在“第一次使用AI”时整体降临的，而是在每一个“这次可以绕过去”的分叉选择里，被一次又一次地结算进个体的行动模式里的。每一次结算，都包含着至少三种条件的共同作用——外部评价的即时反馈、时间压力的紧迫程度、上一次绕过成功留下的神经强化——以及一个逐步固化的自我叙事。这些条件互相咬合，形成了一条自加速的通道：外部越成功，内部越萎缩；内部越萎缩，越依赖外部维持成功。一条互锁的死循环由此结成。

在下文中，我将把这个结算机制拆解为三个互相衔接的动态环节——诱绕、断修、伪稳——但这不是一个可以脱离具体分叉去套用的抽象模型。这三个环节不是三个阶段，不是先诱绕、再断修、最后伪稳的线性进程。它们是同一条自加速通道上三种互相包裹、互相加剧的力：诱绕在每一次分叉处为断修提供条件，断修在每一次挫败中为诱绕加固关口，伪稳在每一个外部肯定的时刻把两者锁得更死。三者不是先后展开的步骤，而是同一次自加速的三种剖面。我之所以分开来讲，只是为了在分析中把它们看清楚——不是为了让它们变成一个诊断流程图。

2.3 诱绕：被短路的锤炼

任何一个真功能——肌肉的爆发力、记忆的提取、从混沌中理出秩序的生成力——都遵循用进废退的铁律。锤炼是所有真功能生成的唯一路径。活壳的致命之处，在于它用极高的效率，诱导宿主逐次跳过锤炼。

2.2节已经展示了一个完整的“逐次绕过”序列。这里不再重复描述，而是追问那个序列中一直被描述、却从未被解释的东西：为什么每一次“绕过”的选择，都会有更强的引力？这不是一个修辞问题，而是一个必须被具体指认的发生学问题。

第一次绕过时，引力来自三重条件的合力：时间压力放大了效率权重，外部评价的历史反馈积累了“绕过更划算”的隐性知识，认知阻力本身的生理成本产生了对“硬扛”的身体性排斥。这三种条件在第一次分叉中合力击败了残余信念。

第二次到第五次绕过时，一种新的条件被加进来了：神经强化。第一次绕过成功后，大脑释放的多巴胺将“绕过→完成任务→获得正面反馈”整条通路标记为有效路径。下一次面对相似阻力时，这条通路会被优先激活——不是在信念层面（“我认为应该绕过”），而是在神经层面（“绕过这条路比硬扛更亮”）。

第六次到第十五次时，又一个条件入场：自我叙事重组。每一次绕过之后，她都在心里对自己说了一次“我不太会写论文”——而不承认的是，这个“不太会”恰恰是她自己在过去半年里每一次选择绕过之后结算出来的功能性退行。但她不知道自己的选择正在结算出这个结果，她以为这就是“我不行”。“我是那种需要工具帮助的人”这个自我认知，在每一次绕过中被反复确认，最终变成了一个稳定的身份叙事。而这个身份叙事，又反过来为下一次绕过提供了最强有力的合法性——“我本来就不会，所以用工具是合理的。我是在弥补自己的短板，而不是在逃避锤炼。”

至此，诱绕完成了从“外部条件驱动”到“内部身份固化”的质变。它不再需要时间压力、外部反馈、同伴效应这些外部条件来推动——它已经被内化为一个自我维持的信念系统：“我天生不行，工具是我的拐杖。”这个信念一旦立住，每一次绕过就不再是一个“被外部引力临时拉过去”的选择，而是一个“与自我认知一致的自治行动”。绕过从有意识的策略变为自动化的惯性，最终固化为神经通路层面上的默认路径：那条通往真功能的路，因为长期无人行走，长满了荒草。

2.4 断修：被熔断的自启信念

诱绕若能被及时察觉和逆转，萎缩仍可挽回。但活壳的第二个动态环节——断修——使得逆转变得极端困难。活壳不仅诱导你绕过功能输出，它还同时在摧毁你“重新启动这一功能”的内在根基：一种我称之为生成性自我效能的深层信念。

生成性自我效能不是对“我能做某件具体的事”的静态自信，而是一种动态的、源自身体深处的、需要被反复验证的信念：我相信，当我把自己丢进一团混沌时，我有能力从这团混沌里，亲手理出秩序，找到出路。这种信念不是被教出来的，它是被无数次的“亲身完成”所喂养长大的。每一次你独立完成一件事——从独自修好一个漏水的水龙头，到独立写出一篇从混沌中理出线索的论文——你的大脑都会收到一条内部的确认信息：“我能行。”这种确认信息不是语言层面的自我暗示，而是由实际行为引发的、在神经层面被编码的自我效能记忆。

阿尔伯特·班杜拉在自我效能感的经典研究中已经确证：自我效能的最强来源是“掌控经验”——即个体亲自完成一项困难任务并成功克服障碍的体验。而活壳所做的，正是系统性地切断了个体获取掌控经验的通道：不是任务变得更难了，而是更难的那部分被外部结构绕过去了；个体看似“完成”了任务，却从未经历过那个“我在阻力中往上顶”的核心时刻。

但班杜拉的框架在这里出现了一个缺口。按照班杜拉的逻辑，每一次“成功使用AI完成任务”的经验，应该转化为对“使用工具”的自我效能——你会越来越相信自己“能用好这个工具”。但活壳揭示的是：这种替代性成功经验不仅没有转化为生成性自我效能，反而在消解它。为什么？

因为掌控经验转化为自我效能的前提，是个体对成功的原因做出内部归因——“是我做到了”。而活壳在每一次完成任务时，都在模糊归因线索。那个高分，究竟是因为我写出了好的文本细读，还是因为AI搭的框架太好了？那个点赞，究竟是因为我的内容有洞察，还是因为算法把我推给了一批刚好会点赞的人？归因模糊，效能就无法被稳固地锚定在任何一极上——既没有锚在“我能用好这个工具”上（因为你知道核心部分不是你做的），也没有锚在“我能自己完成”上（因为你从未真正自己完成过）。效能被悬置在两极之间的真空中，最终从悬置坠入了消解。

它的致命后果，在一个最关键的时刻最为显露：当你决定戒断活壳，重新尝试自己完成一项任务时。你打开空白文档，那个熟悉的阻力感涌上来——它一直都在那里，从来没有消失。但这一次，没有外部引擎可以帮你绕过去了。你坐了片刻，感到

一种从根部升起的恐惧：不是“这很难”，而是“我不知道该怎么办”。那个本应在面对困难时自动启动的“咬牙往上顶”的机制，并没有被激活。你发现，这条通路不仅因为长期未被使用而变得迟钝——它的启动机制本身，已经在无数次的“绕过”中被一个更高效的替代回路给绕过去了。你不是不想往上顶，而是那个“顶”的开关，已经有太多次收到了“不用按”的通知，而在这一轮没有人替它按下时，它自己不响了。

这就是生成性自我效能的熔断——不是功能衰退，而是启动机制被短路。一个陷入断修的人，不是功能上“无能”，而是启动上“瘫痪”。他能清楚地意识到“如果去做，应该这样做”，但他的身体无法接到执行的命令。这种“知道但不行动”的分裂，是断修最精确的体征。

2.5 伪稳：高效如何成为囚笼

功能在被绕过，自启能力在熔断，宿主此时陷入了活壳最具欺骗性的第三个动态环节：伪稳锁死。

活壳替代真功能后，它输出的结果往往在外部评估体系中比真功能更高效、更稳定、更少犯错。这使得依赖活壳运转的个体，从外部看不仅没有“坏死”，反而比那些还坚持用笨拙缓慢的真功能运转的个体，显得更“能干”。所有外部的测量仪表——分数、绩效、产出量、响应速度——都指向同一个结论：这个系统运转良好。

这就是伪稳——一种所有外部指示灯都在闪烁、但内部引擎已经熄火的运行状态。

伪稳不是静止的稳定，而是一个持续加深的旋涡。每一个来自外部的肯定——高分、加薪、点赞——都会在个体内部引发一个互相锁死的双重效应：一方面，活壳的可靠性被再次确认；另一方面，那个曾经怀疑“我是不是太依赖这玩意儿了”的微弱声音，被成功的甜味压下去。外部越成功，内部越萎缩；内部越萎缩，越依赖活壳维持外部成功。一个互锁的死循环由此形成。系统被越来越深地锁死在活壳上，而那个萎缩的真功能，不仅再无恢复之日，甚至被整个系统从记忆中删除了。

至此，三个动态环节走完了一个完整的自加速通道：诱绕让锤炼在每一次分叉处被绕过，断修让重新启动变得不可能，伪稳让整个系统在外部的成功中持续强化这个死循环。每一个环节都在为下一个环节蓄势，整条通道一旦形成便具备了强烈的自加速倾向——不是谁在故意使其恶化，而是结构本身在每一个节点上都给出了“最容易走的路”，而那每一步“最容易走的路”，都在把人进一步推向那个空腔。

2.6 活壳的不可还原定义与三层含义

至此，活壳不再是一个被预先定义的抽象概念，而是一个从具体分叉序列中被逼出来的结构命名。现在可以给出它的定义了。

活壳，不是“外部结构替代内部功能”——那是可以被工程学和病理学已有概念覆盖的描述。活壳切开的，是一个更深的本体论层面：任何因高度适应外部而稳定下来的外部承接性结构，其在成就宿主的同时，都潜在地抽空了那个使它诞生的活的转化力；当这一抽空越过临界线，“成就”与“坏死”便不再是两件事，而是同一过程的两张面孔。

这里需要对“活的转化力”做一个清晰的操作性界定。本文用“活的转化力”指涉主体在面对不确定情境时，能够从内部生成出方向、秩序与意义的能力——它包含三个核心维度：（1）生成方向：在没有外部指令时确定“要往哪里走”的能力；（2）生成秩序：从混沌中理出结构、从碎片中建立连接的能力；（3）生成意义：在挫败与徒劳中维持内在驱动、从自身行动中持续产生“值得做”的理由的能力。这个词中的“活”字仅指“需经亲自锤炼方可生长且不可被外部高效替代的生命功能”——它不预设任何本体论意义上的“本真状态”，而是指向一种必须在反复使用中维持、一旦停止使用便会发生质性坏死的生物性基底。

这个定义包含三个核心构件。其一，悖论性成就：活壳不是失败的结构，而是成功的结构。它之所以成为“壳”，恰恰因为它太成功了——它让宿主在外部评价体系中获得了稳定的、可预期的、高效的表现。成就越大，壳越硬。这是活壳与所有以“失败”或“失调”为起点的既有诊断最根本的分野。其二，假体宿命：任何外部承接性结构，只要它在一段时间内持续有效地承接了主体对世界的回应，它就潜在地

成为假体。这不是因为技术本身有“邪恶属性”，而是因为结构存在本身就有一种惯性引力——它倾向于维持自身、扩展自身、并让宿主绕过那些可能威胁它存在的“笨拙锤炼”。假体不是被设计出来的，是被惯出来的。任何支架，只要拆除的按钮没有被及时按下，就会在引力作用下慢慢滑向假体。这是所有支架的宿命，不是某一类坏支架的特性。其三，坏死不可逆：活壳抽空的那个“活的转化力”，在失去锤炼条件后，不是“休眠”，而是发生质性的、不可逆的坏死——对于已经越过临界线的群体。对于仍处在支架滑向假体过程中、尚未完成结构性消逝的年轻群体，坏死的程度仍是开放问题，交由实证检验。

为防范概念滑动，本文在此明确：活壳在全文中有三层含义——（1）本体论层：悖论性成就结构的共相，一个显露态；（2）动力学层：由诱绕-断修-伪稳三条互相加速的力所结算出的发生过程，一个自加速通道；（3）方法论层：用于诊断上述结构-过程复合体的分析框架。三层不可互相还原，但均以上述定义为共同内核。

三、与近邻的判决性对质

一个理论的成立，不仅在于自身的逻辑自洽，更在于它能清晰地与已有的、试图解释同一组现象的经典判断划清分离线。活壳的不可还原之处，不在于它比既有概念“更精确”——那种在同一个概念光谱上把刀口往左移几厘米的操作，仍然是在旧坐标轴上标新点。真正的不可还原，是在分析推进到某个矛盾的深处时，你发现现有的概念光谱本身都不够用了——你不得不发明一个新的维度。

本节不打算让活壳依次与七个对手对质。许多对手的问题已经在更深的对手那里被覆盖了。我只选两个最强劲、各自代表一条独立论证路径的对手——延伸心灵论（认知科学路径）和斯蒂格勒的无产阶级化理论（技术哲学路径）——以及一位在技术-身体这条线上不可绕过的先行者麦克卢汉，让活壳与他们正面交锋。对质的目的是在对手解释失败的那个具体裂缝处，让活壳的不可还原性自己显影出来。

在进入对质之前，先简要划清活壳与几个常被联想到的概念的边界——这些概念在初稿中有独立对质，此处压缩为快速切割。

与“异化”对比。马克思的“异化”描述的是劳动者与其劳动产品、劳动过程、类本质、他人的分离——你的创造物反过来压迫你。活壳不是分离，而是替代——你的创造物温柔地替你做了你原本需要亲自做的事，而你在这个过程中被系统地绕过了锤炼。异化的主体仍在创造，只是被剥夺了创造的意义；活壳的主体连创造这个动作本身都在萎缩。

与“废弃”对比。废弃描述的是无人使用的退场——一项制度被束之高阁，一项技能被遗忘在角落。活壳描述的则是被替代物高效使用后的退场——功能不是被遗忘的，而是被一个更高效的外部结构温柔地承接走了；而在这种高效使用的过程中，那个被替代的原始功能，因为完全不再被征召，静悄悄地坏死了。前者是弃用，后者是代用。

与“内卷”对比。内卷的主体是在过度竞争中被消耗的——他太累了。活壳的主体是在高效顺滑中被替代的——他太顺了。内卷是“做得太多、累死了”，活壳是“做得太顺、空掉了”。

这三个切割建立了活壳的基本区分面。接下来，是真正的硬仗。

3.1 与麦克卢汉对质：截除是代价，活壳是自噬

马歇尔·麦克卢汉比斯蒂格勒更早地提出了“延伸即截除”——每一种技术的延伸都在人体中造成一次对应的截除。文字延伸了记忆，截除了口语文化的全感官参与；车轮延伸了脚，截除了徒步的耐力；电子媒介延伸了中枢神经系统，截除了个体独处的深度。在麦克卢汉看来，截除是延伸必须付出的代价——你不可能既得到延伸又保留被延伸的原始功能。他的叙述有一种悲观的淡然：截就截吧，这就是媒介的历史。

活壳与麦克卢汉有深层的共振。两者都看见了外部结构对内部功能的置换。但分离线在于：麦克卢汉的截除是延伸的代价——你主动伸出手臂（延伸），手臂末端的某根神经被截断了（截除），代价在你伸出手的那一刻就已完成。活壳发现了比这更黑暗的东西：外部结构不仅在成就宿主的同时截除了某些功能，而且这种截除本身会

因为成就的持续扩大而不断加深——截除不是延伸的一次性代价，而是成就本身的自我强化产物。成就越大，壳越硬；壳越硬，里面越空。

麦克卢汉的截除是一笔一次性付清的代价。活壳的自噬是一笔利滚利的债务——每一期的成功都在为下一期的萎缩增加本金。麦克卢汉没有追问的那个问题是：截除之后，那个被截除的器官，能不能再长出来？活壳的回答是：如果截除是通过“系统地切断锤炼”完成的，那它就不是截除，而是坏死——一旦完成，不可逆转。

判决性预测：如果在一次大规模的“数字戒断”实践中，长期深度依赖数字工具的人能够较快恢复被截除的功能（如独立写作、深度阅读、不依赖外部反馈的内在驱动），那么麦克卢汉的“截除”模型就足以解释问题——截除是可逆的适应。如果戒断后暴露的不是“生疏”而是“空腔”——不是“有点生疏但还能做”，而是“根本不知道从哪里开始做”——那么活壳的坏死诊断就被增强。

3.2 与延伸心灵论对质：工具接管的是手，还是方向？

安迪·克拉克和大卫·查尔莫斯在《延伸心灵》中提出了一个极具挑战性的主张：认知不必被限制在头骨之内，外部工具可以合法地成为认知过程的组成部分。当一个盲人用手杖探路时，手杖已经融入了他的感知系统；当一个数学家在纸上做运算时，纸笔已经融入了他的思维过程。根据这一观点，依赖AI或算法，不一定是“能力的萎缩”，而可能是“认知的延伸”——一种新的、更高效的认知形态正在形成。如果克拉克是对的，那么活壳所诊断的“坏死”，不过是旧范式的认知在向新范式过渡时的不适——就像文字刚发明时，苏格拉底也曾担忧记忆会坏死，但事实证明那是一次认知的跃迁。

这个挑战是致命的。活壳必须清晰地回答：延伸心灵论在什么条件下成立，又在什么条件下不成立？

分离线不在“是否使用了外部工具”，而在工具接管的究竟是认知的哪个层面。延伸心灵论的经典案例——手杖、纸笔、计算器——接管的都是认知的执行层：感觉信息的传导（手杖）、信息的暂存（纸笔）、运算的速度（计算器）。在这些案例

中，认知的生成层——那个人要去哪里（盲人的行走方向）、要写什么（数学家的论证目的）、为什么要做这个运算——仍然由主体从内部生成。手杖没有替盲人决定“我要去哪里”，纸笔没有替数学家决定“我要证明什么”，计算器没有替学生决定“我要解这道题”。延伸是执行层的延伸，生成层没有被接管。

而活壳诊断的对象——AI替你构思文章的结构、算法替你决定下一个该看什么——它们接管的恰恰是生成层：那个“我要说什么”“我要往哪个方向想”“我对什么感兴趣”的核心认知动作。当一个工具不仅替你执行，还替你决定了“你要去哪里”时，它就不再是延伸，而是替代。

延伸心灵论者可能会咬住一个模糊处：AI生成的“论证框架”，在运作机制上完全是执行性的——大语言模型本质上是在做概率预测，不是在做存在性的判断。它没有“想说”什么，只是预测了一个人类读者会认为合理的文本序列。那么，它何以被归类为“生成层替代”？

这正是活壳理论必须钉死的一笔：执行层替代与生成层替代的边界，不在工具的内部运作机制里，而在宿主使用这个工具时的接受姿态里。同一个AI生成的论证框架，如果学生把它当作“一个可以参考的草稿”，自己仍然在其中做二选一、三选一的生成性判断——这个框架放在这里，那里放什么、不放什么，这两个论点谁先谁后，这些仍然是我在拿主意——那么它就在延伸的边界内。它的角色是手杖：给我一个支撑点，但走哪个方向是我。如果学生把它当作“这就是我的论证了，我只用往里面填材料”，那么它已经在替代生成层了——哪怕AI生成它的机制完全是执行性的，它在学生后续写作中的功能已经变成了“要写什么”的生成层替代物。从接受的那一刻起，学生“要说什么”的生成动作就已经被绕过去了。

延伸心灵论的“延伸”前提，恰好是主体在认知过程的生成层仍保持着主动的归因姿态——盲人拿着手杖时，手杖不是替他去他想去的地方，而是帮他更有效地去他想去的地方。当这个归因姿态被系统性地绕过——不是单次，而是成习惯——延伸就变质为替代。延伸心灵论缺少的正是这个区分：同一个外部结构，在同一个体身上的角色，会因为接受姿态的不同而滑落。支架与假体的边界，不在工具身上，在每一次

接受的那个瞬间——在那个瞬间，她是把它当作“我的判断的辅助”，还是“已经替我完成了判断”。林晚从第一次“犹豫”到第十次“不再犹豫”的滑移，正是这个接受姿态从支架向假体的质变。

判决性预测：这个区分给出了一个最强的检验条件。如果在撤除外部工具后，被诊断为“支架性使用”的群体能够较快恢复独立认知功能（因为生成层从未被替代），而被诊断为“假体性接受”的群体暴露出结构性空腔（生成层已被系统性绕过），那么延伸心灵论的通用性就被削弱，活壳的“替代性坏死”被增强。反之，如果两者在戒断后恢复轨迹无显著差异，活壳被推翻。

3.3 与斯蒂格勒对质：被剥夺的知识 vs 被替代的生成力

贝尔纳·斯蒂格勒在《技术与时间》中提出了“无产阶级化”理论：知识的代具化导致个体知识的丧失。当技术装置将知识外化——机器将工匠的手感剥夺，数字平台将消费者的记忆外化——个体失去了知识，变成了某种程度上的“无产者”。这是活壳在欧陆哲学中最直接、最不可绕过的近邻。

现在，我不做概念辨析。我拿一个具体的现场来检：林晚在开题答辩时被导师追问“你论文的核心判断到底是什么”，她无法应答的那个静默。

斯蒂格勒能解释到哪一步？他能解释到“她的知识被外化了”——在过去一年半里，她的论文框架由AI生成，她失去了独立组织论证结构的知识，正如工匠失去了手感。在他的框架里，这是典型的无产阶级化：代具剥夺了知识。所以他的药方是“重新占有”——通过某种新的技术药典，或一种贡献性经济，让她重新接触被外化的知识，恢复独立论证的能力。

但在林晚的静默里，有些东西是斯蒂格勒的解释够不到的。她不是“忘记了论证结构该怎么搭”。如果给她足够的参考资料、给她一天时间、再给她一个AI助手，她可以做得很漂亮——她的壳是完整的。那个静默真正暴露的，不是“知识被夺走了”，而是她内部的某个东西不在了。那个本应在被追问核心判断时自动启动的“往

上顶”的机制——那种在压迫性追问中，从心里自己逼出一句话、哪怕它不漂亮、但它是我自己从混沌里抓出来的——那个开关，在无数次的绕过中已经不响了。

斯蒂格勒能诊断“知识的外化”，但他诊断不了“生成知识的那块肌肉的坏死”。他的“重新占有”药方成立的前提——那个被剥夺的知识生成器官本身依然完好——在林晚的案例中恰好是被质疑的。她需要的不是“还给她知识”，她需要的是被逼着重新长出那个能从内部生成出判断的东西。而那个东西，连她自己都不知道还能不能长出来。

斯蒂格勒的病人被诊断为“他的工具被没收了”——还给他一套工具，他就能重新成为一个有知识的工匠。活壳的诊断是：“他的双手已经不会做工匠的动作了。你给他工具，他也不知道怎么握住它。”这不是在否弃斯蒂格勒——他的“代具化”与活壳的“假体宿命”有深层的共振。但活壳比他多走了一步：它追问的不是“知识去哪里了”，而是“那个曾经产生知识的主体功能，在被代具温柔替掉之后，还能不能再长出来”。斯蒂格勒的答案是可以——否则“重新占有”就毫无意义。活壳的答案是：不一定——取决于那个代具接管的是执行层还是生成层，取决于支架是否在拆除按钮被按下之前滑成了假体。

判决性预测：如果那些通过“重新占有”项目（如去数字化写作工作坊）重新接触被外化知识的个体，能够逐步恢复独立的知识生产能力，那么斯蒂格勒就足够了。如果他们表现出深层的行为瘫痪——知识被还了回来，工具放在了手边，但他们不知道该怎么用——那么，斯蒂格勒的概念光谱就在这里不够用了。活壳所命名的那个“生成器官的坏死”，就从他的解释裂缝处，第一次被逼到了阳光底下。

四、历史的反驳：凭什么这一次是坏死？

任何试图论证数字活壳导致“不可逆坏死”的理论，都必须正面回应一个强大的达尔文式反驳：你所说的“外部替代导致内部萎缩”，不过是人类进化史上的常态。文字被发明时，苏格拉底曾担忧它会损害人的记忆力——他错了，文字非但没有让记忆坏死，反而释放了大脑带宽去发展更高阶的能力。算盘和计算器替代了心算，但它

们让人类能够处理更复杂的数学问题。何以见得这一次的数字代偿——AI替代构思、算法替代方向感——不是新一轮的认知功能外包，而非你所谓的“坏死”？

这条反驳的整个论证结构，都建立在“亲自处理信念”的基石上——它默认为只有被解放的、被释放的、被更高阶的“亲自处理”，才是真功能的生长。但活壳理论恰恰撤销了这一信念本身，因此它能够看见文字与AI之间的质性差异。

关键变量不是“工具的类型”，而是“工具是否系统性切断了锤炼”。文字被发明后，它确实替代了记忆的精确存储功能。但它没有切断锤炼，因为在文字被发明后的数百年里，个体仍有充足的、无可绕过的锤炼条件。在没有搜索引擎和AI辅助的年代，一个人要写出一篇论证，仍然必须亲自承受从混沌里理出秩序的全部压力——文字只是帮助你存储你已经理出的秩序，它没有替你理出秩序。那个“理出秩序”的核心认知动作，仍然在每一次写作中被反复征召、反复锤炼。文字没有切断锤炼，它只是给锤炼的产物提供了一个外部存储器。计算器替代了心算的速度，但它没有切断数学思维的核心锤炼——理解模型、设定变量、解释结果。那个“为什么要做这个运算”的生成层，仍然在主体内部。

今天的AI助手和算法推荐，则在很多场景下直接切入并接替了那个“理出秩序”的动作本身。当一个学生使用AI来生成论文的论证框架时，他跳过的不是“存储记忆”或“运算速度”，而是“从混沌中理出结构”这个核心锤炼环节。这给出了一个更精确的区分判据：如果一个外部工具替代的功能，可以被清晰地界定为“如何做一件事”中的可算法化部分——即可以被提前编程、不需要在情境中即时做出判断的环节——那么它是工具性延伸；如果它替代的功能涉及“这件事要怎么做”中必须在情境里即时生成不可提前编程的判断的部分，或者涉及“为什么要做这件事”“要做什么事本身”，那么它已经进入了主体性替代的疆域。

但必须立即对此判据做一笔关键的限定。执行层替代与生成层替代的边界，不在工具的内部运作机制里，而在宿主使用这个工具时的接受姿态里。同一个AI生成的论证框架，如果学生把它当作“一个可以参考的草稿”，自己仍然在其中做二选一、三选一的生成性判断——那么它仍然在延伸的边界内。如果学生把它当作“这就是我的

论证了，我只用往里面填材料”——那么它已经在替代生成层了，哪怕AI生成它的机制完全是执行层的。活壳的边界，就画在那条姿态的转变线上：从“参考性接受”到“替代性接受”。前者是支架，后者是假体。同一个外部结构在同一体身上的角色滑移，关键变量不是工具本身的能力，而是宿主在每一次面对产出时的归因姿态。

写到这里，必须再对自己追一刀：这个“我判定这一次是坏死”的动作，本身有没有可能是历史上反复上演的“狼来了”故事的最新一个条目？苏格拉底担心文字坏记忆，王阳明担心书籍障良知，卢德主义者担心机器毁手艺——活壳凭什么说自己不是这份名单上的最新一条？

活壳理论对此的回答，不在经验事实层面（因为那是未来的事），而在方法论层面：活壳是第一个为自己的宣告设置了清晰、可被未来检验的证伪条件的同类理论。苏格拉底没有说“如果在文字发明后三代人内，我们仍能观察到不依赖文字的深刻思辨力，则我的担忧被推翻”。活壳做到了这一点——它在最后一章给出了三个明确的证伪条件，其中每一个都可以被未来的人类行为明确地判定为成立或不成立。活壳不要求你相信它——它要求你记住它的证伪条件，然后在二十年、三十年之后，去亲眼看看那片废墟上有没有长出新的森林。一个可以被杀死的宣告，和那些不能被杀死的宣告，在认识论上的地位是不同的。这是活壳对自己最诚实的地方——也是它将“判断”与“预言”划清分离线的理论手术。

五、假体宿命：活壳不可逆的根源

活壳的不可逆性，并非来自外部技术的强大锁定，而是来自它自身内部的结构宿命。这个宿命，可以用一句话概括：任何支架，只要拆除的按钮没有被及时按下，就会在惯性引力下滑向假体；而一旦滑成假体，拆除的按钮就永远失效了。

这一宿命植根于人类生存的根本困境：我们只能在已经形成的结构中应对世界，而这些结构越是成功，就越会取代那个曾经使它们被创造出来的活的转化力本身。这不是某一类坏结构的特性，而是所有结构——无论其设计初衷多么良善——都潜在地

具有的引力方向。支架是为生长而设的临时依靠，但支架一旦立稳，它就开始产生自己的引力场——它让你绕过锤炼后获得的正反馈（更高效、更少挫败、外部评价更高），会诱使你不屑按下那个拆除按钮。而每一次该按却没按的瞬间，都在加固这个引力场，直到最后，支架本身变成了宿主的永久假肢。

这一洞见，逼我们重新审视那个在当代批判话语中几乎被视为自明的“能量隐喻”。这隐喻是这样运作的：将人的内在动力想象为一种可被压抑、疏导、释放的能量。马尔库塞的“压抑性去升华”是其经典表述：势能被虚假满足提前泄掉了。佐博芙的“掠夺”是其变体：能量被盗走了。斯蒂格勒的“无产阶级化”是其欧陆版本：知识被外化了。

这个隐喻看似强大，却含着一个危险且从未被质疑的前提：势能的源头是自足的、不可被消耗殆尽的。它默认那个产生欲望、情感和批判能力的主体内核，就像一口永不枯竭的井，只要不再往里面填石头（压抑）或不再从里面抽水（掠夺），泉水就会重新涌出来。批判理论的全部希望——觉醒、抵抗、夺回、重新占有——都建立在这口井的永不枯竭之上。

活壳撕开了这个隐喻的假面：能量可以从源头处坏死。不是泉眼被堵住了，而是泉眼本身干涸了。那些等着被解放、被唤醒的“内在势能”，在很多长期依赖活壳的个体身上，已经不是一个被压制的地下水库，而是一片沙化的河床。当外部脚手架被撤去的那一刻，暴露出来的不是汹涌的被压抑的激流，而是空洞的寂静。

这里必须立即做出一笔关键的自反，以防范一个易滑的误读：对“沙化的河床”这一意象及其暗示的“从丰沛到干涸”叙事，本文不做价值缅怀。活壳要揭示的不是“曾经有一个好的状态后来坏了”，而是：那个在某些历史条件下被反复征召、反复锤炼而呈现为“丰沛”的生成力形态——那种基于印刷文化、现代教育制度、个体核心劳动形态的特定转化力配置——在新的历史条件下（数字技术、算法环境、外部承接性结构的普遍化），正在被一种新的转化力配置所悬置和覆盖。这个悬置过程本身不是“恶”的，但它有一个结构性的风险：新配置是否切断了旧配置得以自我修复的那个锤炼回路？如果切断了，那么在新配置尚未能自主生成出同等深度的内部转

化力的条件下，个体将面对一个危险的过渡期——旧力已去，新力未生。活壳批判的，就是这个过渡期的结构性风险，而不是对“旧力”的怀旧式追挽。

与这一自反相配套，本文也需对“亲自处理信念”做一笔清晰的切割。本文对“活的转化力”的指认是发生学的——它只能以“亲自承受某种特定类型的阻力并在此阻力中反复锤炼”的方式生成。在当前人类的生物条件和神经可塑性机制下，这是一个发生学事实，不是一个价值判断。如果未来的技术发展使得生成性自我效能可以通过非锤炼途径建立（例如神经接口技术直接刺激前额叶皮层相关回路），那么活壳的坏死诊断就需要被修正。本文不预设“亲自”具有内在的道德优越性——它只是在当前的生物条件下，恰好构成某种特定能力的必要条件。这一区分必须在全文贯穿始终，否则“坏死”“空腔”等措辞将不可避免地携带一种未被承认的怀旧的道德气味。

于是，我们终于抵达了活壳最深的那道裂缝：人类迄今为止的整个价值秩序——特别是其教育与成长叙事——都建立在“亲自处理信念”的基石上。但那个曾经铸就了这一信念的社会条件，正在被活壳无声地侵蚀。当“亲自去做一件事”的机会成本越来越高——因为每一次笨拙的尝试，都可以被流畅的外部替代所绕过；当那些支撑“亲自处理”的外部结构——师徒制、需要长期浸染的学徒暗夜——正在被效率至上的社会机器系统性地拆除；当“亲自去磕”不再被看作成长的必经之路，而被看作可以被优化的低效冗余——此时，活壳就不仅是个体的病理，它正在成为社会结构的默认运行模式。

一批又一批的年轻人，正被温柔地、高效地、用最好的名义，从“亲自处理信念”的基石上连根拔起。他们被拔走的时候，甚至不觉得痛——因为活壳给了他们顺滑、高效、不出错的替代体验。直到有一天，他们面对生命中第一件没有模板、没有参考、没有任何外部引擎可以替他们从混沌中生成出秩序的真正创造时，才会在那种从根部升起的瘫痪感中，第一次触摸到那个空腔。

六、从支架滑向假体：一个完整的发生学案例

活壳的三条自加速力——诱绕、断修、伪稳——在第二章中是在微观分叉序列里被追溯的，它们的共同作用构成了一个自加速通道。本章将把这个通道完整地铺在一个具体的人身上：从她第一次接触数字活壳，到支架向假体的滑移，再到某个危机时刻空腔的暴露，最后到外部仪表的轰鸣与内部引擎熄灭的同步被看见。这个案例不是要“证明”活壳理论——一个单一个案证明不了任何理论。它的功能是发生学的：让那个在抽象分析中容易被当成“诊断框架”去套的东西，被一次具体的生命轨迹逼回它的发生现场，让读者自己亲眼看见活壳是如何在一次次选择中被结算出来的。

6.1 案例：林晚（化名）的发生学轨迹

林晚，25岁，中国某高校中文系硕士。她本科阶段以独立写作出众，曾获得校级散文奖。2023年初，她在同学推荐下开始使用一款大型语言模型辅助论文写作。

在林晚滑入活壳之前，有三重条件使她比同龄人更容易陷入这个通道。其一，她对“写好论文”的在乎程度极高——本科的散文获奖让她在进入研究生阶段时，将学术写作能力锚定为自己身份的核心构件。这意味着面对每一次写作任务，外部评价在她心中的权重都远高于平均水平。其二，她所在的研究生培养环境本身，已经是一个更宏观的活壳体系——标准化评价模板、按时毕业的倒计时压力、导师在有限的指导时间内不得不更看重“论文能不能通过”而非“学生是否在真正生成判断”。这个大活壳给她的小活壳提供了土壤。其三，她第一次面对AI辅助时的那个关键时刻，没有一个能逼她“自己扛过去”的外部导师在场。她是在独处中做出那个第一次绕过选择的——没有旁观者，没有叫停者，没有一个来自外部的声音说“你应该自己试试”。

这三重条件的共同作用，使得林晚在面对活壳的第一个分叉点时，被拉向“绕过”一端的引力格外强。以下序列是她的轨迹。

2023年1月：第一次使用。起初，她只是让AI帮忙润色几个句子。这是支架式的使用：她在修改时仍在思考“为什么这个词比那个词更恰当”。锤炼仍在发生。

2023年3月：第一次绕过。春季学期的一篇课程论文。面对一个艰涩的文本分析段落，她在空白文档前坐了四十分钟，感到强烈的思维阻力。她想起同学说“直接丢给AI生成框架再改”。她犹豫了几分钟。在这个分叉点，三重条件同时在场：截止日期在逼近，上一门课自己硬写得了低分而用AI的同学得了高分，那个阻力区本身的生理性不适让她本能地想要回避。三重力叠加，她选择了“绕过”。AI在三秒内生成一个论证框架，她检视、修改、补充，提交。成绩87分——比她上一篇完全自己硬写的论文高了五分。第一次绕过成功。

2023年秋季：第五到第十五次绕过。她在这个学期写了四篇课程论文，其中三篇都使用了同样的工作流：先看AI能生成什么框架，再在上面修改补充。“先自己试试”的动作已经被从流程中删除。这不是因为她在信念层面做出了“我要永远依赖它”的决定，而是因为每一次绕过后，正反馈——高分、顺滑、少受罪——在神经层面反复强化了那条替代通路。与此同时，一种归因模糊正在她体内累积：那个高分，究竟是因为她写出了好的文本细读，还是因为AI搭的框架太好了？她说不清楚。说不清楚，效能就无法被锚定。效能被悬置在两极之间的真空中。

2024年2月：空腔的第一次暴露。她在一个没有网络的学术沙龙上，被要求即席发言十分钟，主题是她正在做的研究。她站起来，说了不到两分钟，就开始语无伦次——她能够感觉到那个曾经在本科阶段反复训练过的“从混沌中理出结构”的能力此刻想要启动，但那条神经通路已经被无数次绕过盖上了厚厚的一层替代性习惯。她慌乱地放弃了发言。那一刻，她第一次触摸到自己的空腔。

2024年春：第一次判决性对比——撤除工具。研二下学期，她需要独立完成一篇一万五千字的学术论文作为硕士论文开题预备。这一次，她试图减少对AI的依赖——导师建议她“还是要独立写”。她打开空白文档，试图自己搭框架。那个熟悉的思维阻力涌上来。这是第一次“撤除外壳”的场景：工具被撤除了。但出现的不是“生疏”——不是“有点生疏但练几天就好”。出现的是瘫痪：她整整坐了一个下午，删了写，写了删，最终一个字也没留下。她后来说：“我不是不会写。我知道如果让AI生成框架，我可以在上面改得很漂亮。但我自己从那片空白里开始的时候，我的整个身体都在抗拒。好像有一条路，我清楚地记得自己曾经走过，但现在那条路上全是

雾，我看不见入口在哪里。”这一次撤除暴露的是生成层的瘫痪——空腔的剖切面

A：从空白中生成结构的能力，已经在无数次的绕过中被拆解了。

2024年秋季：第二次判决性对比——撤除模板。硕士学位论文开题答辩。她提交的开题报告中，论证框架部分主要由AI辅助生成，她在此基础上做了详细的文本细读补充。答辩时，评审老师对她的细读部分给予了高度评价，但追问了一个她从未准备过的问题：“你论文的核心判断到底是什么？三个分论点之间的关系，是你自己想的，还是你只是按一个逻辑模板排列了三条材料？”她愣住了。她发现，“按逻辑模板排列三条材料”是对她过去一年半工作流的精确描述——而她在此前从未意识到这有什么问题。她可以流利地复述报告里的每一个细节，但她无法在教授的要求下，站到更高的一个层级，用一句话说清“你到底想表达什么”。那个在答辩圆桌旁不应答的静默，就是空腔的剖切面B——这一次撤除的不是AI，而是“模板”本身。在第一个撤除中，AI被拿走了，她面对的是“空白文档”。在第二个撤除中，教授恰恰要她站到模板之上的那个位置——去命名那个本该是她自己生成的核心判断。而这种“站到模板之上”的能力，在她过去一年半的使用中，因为从未被征召过，已经被模板本身给替换掉了。

两个撤除暴露的是同一个空腔的两个不同剖面。切面A：没有外部引擎时，自己从空白中生不出结构。切面B：有外部引擎搭好结构时，自己站不到结构之上去说出那个“你的判断”。A是生成层的瘫痪，B是被模板替代的归因主体——她不知道自己到底判断了什么，因为那个判断动作从第一次绕过开始就没有被她亲自做过。

案例的理论回收。林晚的轨迹完整映射了活壳的三条自加速力。诱绕：从第一次绕过时的“犹豫几分钟”，到第七八次绕过时“不再犹豫”，再到第十次之后“AI变成默认起点”——每一轮绕过留下的正反馈（高分、顺滑），都在神经层面强化了替代通路，在自我叙事层面加固了“我是那种需要工具帮助的人”的自我认定。断修：每一次成功使用AI完成任务的体验，都因为归因模糊而无法转化为对任何一极的稳定效能感——既没有锚在“我能用好AI”上，也没有锚在“我能自己完成”上——效能被悬置在两极之间的真空中，最终从悬置坠入了消解。伪稳：从2023年春到2024年秋的一整年半里，她的外部评价体系——课程分数、导师的初步肯定、开题报告的文

本完整度——始终在给出正向信号，直到2024年秋的那个答辩现场，当教授追问的恰好是被模板替代了的那个生成层时，外部仪表才第一次和内部引擎的熄火同时被暴露在光天化日之下。

6.2 边界条件

林晚的案例典型，但它的典型性不能用来论证活壳的必然性。理论如果不能解释反例——那些深度使用工具但生成能力依然很强的人——就还没吃透自己的边界。

以下勾勒一个简要的对比剖面。一位同龄、同专业的男性研究生，同样从2023年初开始深度使用AI辅助写作。但他的轨迹与林晚不同。一年半后，他在类似的开题答辩中，能够流畅地说出自己的核心判断，而且在被追问时能够即时调整论证结构——不是因为反应快，而是因为他每一次使用AI生成的框架时，始终保持着一种“不接受的接受”：他每次都把AI的框架当作一个需要被他推翻、重排、或者至少被他二选一的草稿，而不是当作“我的论证”。他从未删掉自己亲自面对空白文档的那个环节——哪怕只是十分钟，哪怕最后写出的是片段，但他每次都会在那个阻力区里先待一会儿，然后再去使用工具。在他身上，支架从未滑成假体，锤炼从未被系统性绕过。

这个对比案例的功能不是“证明”什么——它是帮我们看清活壳的边界条件：活壳的发生不取决于“是否使用了工具”，而取决于“在使用工具的过程中，生成层的锤炼是否被系统性绕过”。保护生成层的那个变量，不是反技术的禁欲，而是一种特定的接受姿态：把工具的产出当作需要被我推翻、重排、裁决的候选，而不是对我的判断的替代。支架与假体的分界线，就画在每一次接受的那个瞬间。

七、活壳的操作接口

一个理论若“说得很对但别人用不了”，就还只是一件概念假体。本节为活壳装上可操作接口，让一个研究者能在十分钟内用它去分析一个新案例。

7.1 最小分析步骤：四步诊断法

当你面对任何一个表现优异却“不对劲”的系统——一个高分学生、一个高效团队、一个高活跃度的用户社区——你可以按以下四步启动活壳诊断。每一步都配有操作化判据，确保诊断可被复现。

步骤一：定位成就。先不批判。诚实地问：这个系统在外评价体系中，最被肯定的成就具体是什么？是论文的高分？是流量的增长？是绩效的稳定？必须指认到可观测的指标上。操作化判据：该成就是否可以在不依赖任何解释性话语的情况下，被一个外部观察者独立核实？如果不能，说明成就的定位还不够具体。

步骤二：追踪替代。这个成就是如何被达成的？在从混沌到秩序、从空白到产出的整个过程中，是否有某个关键环节——那个原本需要主体亲自去承受不确定性、反复失败、在压力中理出线索的环节——被一个外部结构（模板、算法、制度流程、他人代劳）高效地承接了？这里需要的是发生学的还原：不是问“它用了什么工具”，而是问“在哪个时刻，主体本该亲自承受的压力被绕过去了”。操作化判据：在从开始到完成的整个行动链条中，找出至少一个时刻——在这个时刻，如果撤除外部结构，主体是否还能独立完成该环节？如果主体完全不能辨认出“如果工具不在我该怎么办”，则替代已发生。

步骤三：检测空腔——临界线判据。这是整个诊断最关键的一步。如果暂时拆除那个外部结构（让人离开模板、让平台停止推送、让制度撤销支撑），在同等困难的任务面前，主体表现出的是行为性的“生疏”，还是结构性的“瘫痪”？二者的操作化区分如下：（1）如果是生疏，主体会表现出：能识别任务的核心要求，能说出“如果给我时间练习，我应该可以”，在尝试时表现出逐步改善的趋势，挫败感在尝试初期较高但随着时间的推移而下降。（2）如果是瘫痪，主体会表现出：无法识别任务的入口（“我不知道从哪里开始”），无法调用曾经存在过的内部操作图式（不是“手生”而是“连图式都没了”），在尝试时没有改善趋势——挫败感随时间推移而上升而不是下降。临界线就画在这两种状态之间：当撤除外部结构后，主体不

是感到“困难但可启动”，而是感到“根本不知道从哪里启动”——此时，支架已经滑成了假体，坏死已经完成。

步骤四：判断伪稳。回到外部评价体系。这个系统目前的优异表现，是否依赖于那个外部结构的持续运转？如果是，那么当那个外部结构一旦撤去，系统是否会暴露出一个内部空腔，使得其外部优异的假面瞬间崩塌？操作化判据：做一个思想实验——如果外部结构在明天被永久撤除，主体为了维持当前表现水平，需要多长时间的恢复期？如果恢复期被预估为“不知道，可能永远恢复不了”，则为伪稳；如果预估为一段可标明的时间，则为工具性依赖。

7.2 中间层操作化：编码判据

“四步诊断法”提供的是宏观流程，但在实际分析文本、访谈记录或行为日志时，研究者需要一个更具体的中间层编码判据，以将“诱绕”“断修”“伪稳”三个动态环节从定性叙述转化为可复现的分析动作。以下编码表提供了一套初步的操作化方案，供研究者在使用时根据具体情境调校。

诱绕检测。 编码对象：个体的工具使用自述或行为日志。编码指标：（1）从“先自己尝试再使用工具”到“先使用工具再修改”的次序转变——如果观察期内自述首次使用时为“先自己尝试”，第N次后转为“先使用工具”，标记为诱绕启动。（2）面对任务阻力时“犹豫”时长的缩短——如果首次描述“犹豫几分钟”，后续描述为“不再犹豫”或“直接打开AI”，标记为诱绕加速。（3）归因语言中自我叙事的变化——如果早期自述为“这个我做不好，需要工具帮助”（情境性归因），后期自述为“我就是不太会写论文的人”（特质性归因），标记为诱绕固化。

断修检测。 编码对象：个体在完成任务后的口头或书面反思。编码指标：（1）归因模糊的存在——如果反思中同时包含“AI帮了很大的忙”和“我也做了不少补充”，且未对成功的原因进行明确的内外归因切分，标记为归因模糊。（2）效能语言的悬置——如果个体在反思中使用“说不好是谁的功劳”“不知道如果自己写会怎样”等表述，标记为效能悬置。（3）启动瘫痪——如果在工具被撤除后，个体在尝试

独立完成时报告“不知道从哪里开始”“有种从根部升起的恐惧感”，而非“有点手生但应该能练回来”，标记为断修完成。

伪稳检测。编码对象：外部评价指标（分数、绩效、点赞量等）与内部状态的对比。编码指标：（1）外部指标的持续正向与内部状态的持续恶化之间的背离——如果不能同时获得内外部两套数据，可通过个体自述中的“一切看起来都不错，但我自己知道不对”来识别。（2）对工具撤除的灾难性预期——如果个体在被问及“如果这个工具明天永久下线你会怎样”时，回答中包含“我就完蛋了”“我会被淘汰”等不可替代性表述，而非“需要一段时间适应”，标记为伪稳。

7.3 外推接口：跨域的活壳形态

活壳不仅适用于数字语境。一旦掌握了“成就本身潜在地抽空转化力”这一辨别维度，它就可以被外推到更广的领域。但必须申明：以下外推属于类比论证，它对活壳核心框架的跨域可迁移性提出的是假设而非证明。检验这一可迁移性，需要在这些领域分别进行独立的发生学分析。此处仅勾勒同构形态，供后续研究作为起点。

在教育中，当标准化的教学流程、现成的知识模板、高利害考试的答题框架构成了学生的活壳时，学生拿到的漂亮分数，是壳的分泌物。他们脑中那片应该长着思想肌肉的地方，可能是空腔。这个问题只有在他们面对第一件没有模板的真实创造时，才会被检验。

在组织管理中，当一个企业依赖一整套成熟的流程、KPI体系、标准操作程序来维持高效运转时，它可能在成就自身的同时，系统性地诱绕了组织成员在混沌中独立判断、在不确定中做出原始决策的能力。这个企业越成功，它的活壳就越硬；活壳越硬，它面对颠覆性创新时的内部抵抗力就越弱——因为颠覆性创新需要的恰恰是被活壳抽空的那种“从空白中生成新方向”的原始能力。但需要注意：组织是一个与个体不同的分析单元，组织“生成力”的坏死是否遵循与个体相同的机制、是否同样不可逆，在此处仍是开放问题。

在以下生态学的例子中，本文进入的是隐喻性类比——它跨越了人类主体性的边界，将活壳的逻辑投射到非人类的生态系统上，意在揭示一种同构的模式，而非宣称活壳在生态系统中具备与人类领域相同的本体论地位。现代单一作物农业依赖高产品种和化肥滴灌系统来实现惊人的产量。它替代了那片土地在漫长生态进化中形成的、能够自维持、抗灾害、不断生长的复杂土壤微生物群落。在化肥的支撑下，产量报表漂亮了数十年。但土壤的真功能在持续坏死：有机质枯竭，微生物群落崩溃。这个被架在高产上的系统进入完美的伪稳状态。直到某一天，能源危机导致化肥断供——活壳撤去，那片已经沙化的死土，才在阳光下发出一声不存在的恸哭。

这些在不同尺度上几乎同构的形态，都在诉说着同一件事：当我们用一架漂亮的外部结构，接替了那个需要在漫长岁月中被反复逼到极限才能长出的内部转化力时，我们得到的短期成就，可能正在以一种不可逆的方式，从我们的底层拿走一些我们甚至不知道该如何命名的东西。

八、自反：理论自身的活壳检测

现在，必须用活壳这把解剖刀，反过来切向它自身。

8.1 批判理论与活壳

这种解剖的行为本身，是否也是一件概念上的活壳？当我们用一个精密的理论框架，去识别、命名并批判那些使人类生命萎缩的外部结构时，这种批判的动作，是否也同样替代了那种更笨拙、更缓慢、更沉默的自我修行？它给了我们一种“我正在深刻思考”的假劲，而可能同时荒废了那个单纯的、不做任何理论反射的“沉浸于生命中”的真功能？

对此，本文的回答是肯定而诚实的：哲学与批判理论，如果其终点仅仅是让信奉者获得了一种“我看透了”的智力优越感，而并未转化为一种对自身生活方式的任何微小修正，那么它同样是一件漂亮的活壳。我们用批判活壳的话语来制作自己的活壳，这是所有批判理论都无法逃脱的最大反讽。读过这篇文章后，那种“终于有人把

这种感觉说清楚了”的通透感，本身就可能是一种极高效的替代性满足——它让你在不必改变任何生活行为的情况下，就获得了一种“我已经在对抗活壳”的假劲。

这件活壳要证明自己不是新的萎缩加速器，唯一的方式，就是它必须为自身设定明确的、可被现实事件杀死的条件。一个不能被杀死的东西，没有生命。

8.2 证伪条件

证伪条件的设置借鉴了科学哲学中可证伪性的一般观念，但并不要求逻辑实证主义式的严格判决性实验，而是为理论提供明确的可检验判据，使活壳成为一个可被未来经验修正的框架。

条件一：代际生命力检验。如果有一代人，在从小深度使用各类数字活壳（AI辅助、算法推荐、高度外部驱动的环境）成长起来之后，当他们步入中年，面对生命中真正的、无法被任何模板或算法解决的深渊般的困境时——比如重疾、至亲丧亡、深刻的信仰危机——他们之中有可观比例的人，能够在没有外部支架的情况下，不瘫痪、不崩溃，而是被困境逼出了一种前几代所未知的、崭新的、坚韧的深海生命力，从这片活壳的废墟里重新长出真正的、属于自己的意义底盘。如果这种生命力不是个例，而是一种可观测的代际趋势，那么，“活壳导致不可逆坏死”这一核心论断，就因这一次伟大的物种进化而被推翻。

条件二：戒断恢复检验。如果在大规模的实证研究中，能够观测到深度依赖数字活壳的群体在长期戒断后（例如一年以上不使用任何AI辅助、算法推荐、无限信息流），其原本被认为已经坏死的生成功能——独立组织原创论证的能力、在没有外部反馈时依然保持内在驱动的能力、在没有模板时从混沌中理出秩序的能力——出现了显著的、可持续的恢复，并且恢复程度与戒断时间呈正相关，那么，“坏死不可逆”这一论断就被证伪。它将被降格为“功能暂时性退化”。需注意：此条件与林晚案例中“空腔暴露”的定性并不矛盾——林晚处于支架滑向假体的过程中，尚未被断定为“已越过临界线的不可逆坏死”。条件二的存在，恰恰是为了检验“临界线”是否有回退可能，而非事先宣布它为不可回退。

条件三：架构修正检验。 如果未来出现了一种全新设计的数字环境——它依然提供高效的工具，但系统性内置了“锤炼保护机制”（例如，在提供模板时强制用户先自己尝试；在推送算法中留出不被填充的空白时间；在AI辅助中设置“拆除节点”——在用户连续使用达到某个次数后主动提示戒断）——并且在这种环境中成长起来的一代人，成年后并未出现大规模的生成性自我效能衰退，那么，“活壳是当前数字架构的必然产物”这一诊断就被削弱。活壳将被证明是一种可被架构修正的问题，而非结构的宿命。

这些条件的共同特征，是它们都可以被未来的现实世界明确地判定为“成立”或“不成立”。它们给了这个理论一个可以被杀死的机会。而一个可以被杀死的理论，才是一个有生命的理论——它不是一件漂亮但已经坏死的活壳，而是一件仍在呼吸、仍在等待被现实检验的真东西。

8.3 出路不在回归，而在重认

最后，本文必须处理一个潜藏在全文批判中的危险乡愁。文中对“活的转化力”的反复指涉，是否暗示了一个没有被活壳污染的、纯净的“本真状态”？是否意味着我们应该回到某种前数字时代的、手工匠人式的浪漫过去？

必须明确：不存在那样一个等待被回归的“本真状态”。我们关于“本真性”的渴望本身，就是被特定历史条件发生出来的结构——它是浪漫主义对工业革命的回应，是存在主义对战后虚无的回应。即使是本文反复论及的“活的转化力”——那个从混沌中生成秩序的能力——也不是人的永恒本质。它是在特定历史条件（印刷文化、现代教育制度、以个体为核心的劳动形态）下被反复征召、反复锤炼，从而被强化和维持的生成结构。当那些历史条件发生变化，这个生成结构本身也可能发生形态转换——活壳批判的终点不该是“捍卫这个形态不被侵蚀”，而应是揭示侵蚀正在发生，并公开它的条件，让每一个面对它的人能够在清醒认知的基础上，自主地决定如何处置它。

出路不是“回归真我”，而是：在承认坏死可能已经完成的前提下，接受自己只能从这片空腔之上重新开始——不在于“长出”什么，而在于“怎么长”本身也在这

次空腔的逼迫下第一次被发明。那个新的生成力长什么样，本文不知道。本文的责任仅仅是：把那个正在我们体内发生的坏死，逼到阳光底下，让它不再有藏身之处，并给它装上可以被杀死的条件。至于在那片被暴露出空腔之后，能不能长出前所未知的新东西——那是属于这个时代每一个拒绝被彻底替代的人，自己的事了。

活壳是这样一词：它逼你看见，那些让你活得顺的，也正在让你空掉——而且这两件事的根，是同一条。

证伪条件（用最简单的话）：如果有一代人，在从小深度依赖数字活壳长大之后，步入中年面对真正的深渊时，不仅没有瘫痪，反而被逼出了一种崭新的、前几代所未知的深海生命力，从活壳的废墟里重新长出了自己的意义底盘——而且这不是个例，是一种可观测的代际趋势——那么，活壳的“不可逆坏死”论断就被推翻。

注释

[1] 此处对“活”字的生物学基底说明，意在防范将“活的转化力”误读为某种浪漫主义式的本真自我或形而上学实体。它的存在和消退均可在功能层面被观察和描述，无需预设一个不可分析的内核。本文对“亲自处理”的反复强调，不是对其价值体系的怀旧式重申，而是因为在当前人类生物条件下，某些核心转化力的生成恰好需要一个“亲自承受某种特定类型的阻力”的锤炼过程。如果未来人类的生物基底发生改变，活壳的坏死诊断可能需要被修正。

[2] 班杜拉将自我效能的来源归纳为掌握经验、替代经验、言语说服和生理情绪状态四种，其中掌握经验被证明是最强烈、最持久的来源。本文提出的“生成性自我效能”在来源层面与之重叠，但更强调其在面对不确定性混沌情境时的生成维度，这与班杜拉最初在具体任务情境下的自我效能概念在应用范围上有所不同。

[3] 延伸心灵论的经典论证（Clark & Chalmers, 1998）的核心是，当外部工具在认知过程中扮演了与内部认知过程相同的功能角色时，它就应该被视为认知过程本身的组成部分。本文接受这一论证的形式正确性，但指出认知延伸的合法边界在于被延伸的部分必须是可算法化、无需情境判断的执行性功能。将生成层功能也纳入延伸范围，是对延伸心灵论的过度扩张，而非其逻辑必然。

[4] 此处对文字与口语文化关系的理解，受到沃尔特·翁（Walter Ong）口语文化与书面文化研究的启发。翁的主要关切在于意识结构的转变，本文则将分析推进到锤炼机制是否被切断的发生学层面。

[5] 本案例基于真实人物经历，但主人公姓名、院校等识别信息已做匿名化处理。对话与心理描写根据当事人多次访谈整理，并为了凸显发生学机制的典型性做了必要的合成与简化。本案例旨在展示活壳机制的微观发生序列，属于思想拓展性质的例示，不作为经过同行评议的实证数据使用。

[6] 生态学类比属于隐喻性类比，旨在揭示一种跨域的同构模式，而非宣称非人类系统中存在与人类主体性相同的本体论机制。其功能是启发性的，不是论证性的。

[7] 证伪条件的设置借鉴了科学哲学中可证伪性的一般观念，但本文并不要求逻辑实证主义式的严格判决性实验，而是为理论的可检验性提供清晰的操作判据，使活壳成为一个可被未来经验修正的框架。

参考文献

- Bandura, A. (1997). Self-Efficacy: The Exercise of Control. W. H. Freeman. (Pp. 80–86)
- Bauman, Z. (2000). Liquid Modernity. Polity Press.
- Bourdieu, P. (1979). Distinction: A Social Critique of the Judgement of Taste. Routledge.
- Clark, A., & Chalmers, D. (1998). The Extended Mind. Analysis, 58(1), 7–19.
- Han, B.-C. (2010). Müdigkeitsgesellschaft. Matthes & Seitz. (Gewalt der Positivität)
- Marcuse, H. (1964). One-Dimensional Man. Beacon Press.
- Marx, K. (1988). Economic and Philosophic Manuscripts of 1844 (M. Milligan, Trans.). Prometheus Books. (Original work published 1844)
- Ong, W. J. (1982). Orality and Literacy: The Technologizing of the Word. Methuen.
- Stiegler, B. (1998). Technics and Time, 1: The Fault of Epimetheus (R. Beardsworth & G. Collins, Trans.). Stanford University Press. (Original work published 1994)
- Zuboff, S. (2019). The Age of Surveillance Capitalism. PublicAffairs.

附论零·本组十篇的分工（编辑增补）

说明：本组十篇出自三次连跑（项飙「附近」／鲍曼液态现代性／伊洛思情感资本主义），彼此距离很近。以下分工地图十篇同置一份，并附上「这份分工在什么条件下应当被推翻」——这比事后挑几篇发要诚实。

十篇不是同一判断的十种说法，它们各自占住的是一个不同的自变量。按自变量归并，落在四个面上：

A 认识的边界：守护者／陪伴者知不知道正在发生什么。《喂养附近》问的是守护者的「可言说性自信」；《从守护到长陪》问的是陪伴者在自己也失去方向感时能否继续在场。两篇共享的判断是：不知道，不是待补的缺陷，而是必须被结构性地容纳的条件。

B 发生的条件：什么情境逼出不可代理的动作。《被裁决》给出四重条件（悬空·无代理·不可撤回·他者在场）；《关系的重力》指认关系可能死亡的不确定性是意义得以持续的前提。两篇共享的判断是：不确定性不是意义的障碍，是意义的燃料。

C 沉积与抽空：能力怎么积起来，又怎么被掏空。《被需求反射接生的人》追自我叙事在平台反馈中被重组；《底劲》追前意志层的身体厚度如何积蓄与中断；《活壳》追外部适应在成就自身的同时抽空内在转化力。

D 时序的替代：外部裁定发生在内部生成之前。《关系的腹语之死》《感受生成过程的算法替代》《当亲密不再亲临》共享同一个时序判断——不是既有的感受被篡夺，而是感受还没从互动内部长出来，就被一个外部的、可计量的裁定跳过。三篇的分工是：一篇追关系内不可译层的消失，一篇追替代的运作机制与边界，一篇追替代之后主体侧留下的存在论沉寂。

这份分工的有效反驳条件。若去掉上述自变量之后，其中两篇在同一情境下仍然做出同样的预测，那么它们就应当被合并，而不是各自成篇。

最需要交代的是 D 组内部与 C 组的《活壳》之间。可裁决的分离点是这一格：设想一个人在完全没有任何外部裁定介入的条件下长期生活（无算法、无测评、无第三方解读），但同样长期不必亲自处理任何有阻力的材料。D 组预测他的关系意义不会沉寂（因为时序未被替换），《活壳》预测他仍会坏死（因为抽空来自适应本身，不需要外部裁定）。**这一格是《活壳》不能被并入 D 组的唯一理由。**若实证上这一格不存在，《活壳》应并入 D 组，本组应为九篇。

附论一·最近邻补切与可裁决预测（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、却与核心命名距离最近的占位者，并给出可裁决的分离预测。

增补一 与克拉克「延展心智」的正面对质

本文通篇使用"支架"与"假体"这组隐喻，却没有与提出这组隐喻的理论正面交手：安迪·克拉克与查尔莫斯的**延展心智**论题——认知过程可以合法地延展到颅外的工具与环境中，笔记本对记忆受损者而言就是记忆本身，不存在一个必须被守在皮肤之内的"真正的"认知。若延展心智成立，本文所诊断的"坏死"就只是一次正常的认知边界外移，根本不是损伤。

这是本文必须接住的最强反驳，分离线要切在**被延展的是什么**。延展心智论题处理的是功能的实现位置：一个认知功能由脑内还是脑外的载体承担，在功能主义看来无关紧要，只要耦合稳定、可及、可信。本文关心的不是功能实现在哪里，而是**那个生成新功能的转化力**——延展心智可以解释"我不再记电话号码而记在手机里"为何无损，却无法解释"我不再从混沌中理出秩序"为何有损，因为后者不是某个功能的实现位置改变，而是产生新功能的能力本身停止运转。

可裁决的分离点：给受试者一个**其外部载体无法覆盖**的新问题（没有先例、没有模板、工具给不出答案）。延展心智预测长期使用外部载体的人在此不劣于对照组——因为耦合只是扩展了资源，没有削弱本体；本文预测他显著更差，且差距不随时间练习而缩小。若测得无差异，则活壳可还原为延展心智的一次悲观读法，本文的坏死判断应被撤回。

增补二 站内划界：与本栏《改不动的机器》及高鹏《活性塌缩》《秩序自噬》

"活壳"这个命名在本站内部有三个近亲，必须自己切开自己，否则是最典型的自撞。

与本栏《改不动的机器·自噬性稳态》：那一篇的载体是制度，机制是改变的努力被系统翻译为加固自身的养分，观测点在改革动能的去向。本篇的载体是个体的生

存结构，机制是适应本身抽空转化力，观测点在成就与内在活力的背离。两者的分离判准：自噬性稳态要求存在一个"元标准作为最终仲裁者"的制度条件，抽掉它机制即不成立；活壳不需要任何仲裁者，一个完全没有外部评价的人同样可以在适应中被抽空。

与高鹏《活性塌缩》（本站学员专栏）：那一篇诊断的是司法制度耗竭自身存活条件，观测点在需求侧——承受者停止认领。活壳的观测点在供给侧的内部——结构自身的转化力停转，与外界是否认领无关。**与高鹏《秩序自噬》**：那一篇的自变量是"自运转程度"这一结构条件，是秩序侧因成功而反噬；活壳的自变量是适应深度，且它主张的不可逆性比秩序自噬更强——秩序自噬允许通过修复发生条件回转，活壳主张越过临界线后不可逆。

可裁决的分离：若在同一对象上观察到"外部评价撤除后转化力回升"，则活壳应并入秩序自噬那一族（可修复的成功代价）；若观察到评价撤除后仍不回升、且回升与否只由适应深度预测，则活壳的独立位置成立。这一条同时也是本篇第一条证伪条件的加强版。