

谁在掌尺？——AI 时代经济学中主体尺度再生产的断裂与拮抗

作者 鲍锦朝 · SDE 学员 · 约 1.7 万字 · 发表于2026年7月24日

摘要 · ABSTRACT

AI 对经济系统的深度嵌入，正在引发一记比“主体被替代”更隐蔽的位移。本文诊断的核心对象并非技能的贬值或判断力的磨损，而是一个更为底层的再生产过程的悬停——经济主体用以亲手重制自身价值刻度的那个完整肉身回路（本文称之为“掌尺回路”），因系统性缺乏“亲自承担不可逆后果”这一启动条件，而进入传导阻尼状态。既有进路——工具论、成本论、主体性折旧论——虽层层深入，却共享一个未经审查的预设：尺度是给定的，或可在不亲自重制的条件下被守护。本文的核心工作分三步展开：首先，通过对一个资本配置决策在“亲自下注”与“外包给 AI”两种模式下神经级因果机制的对比推演，从后果承担、躯体标记写入与前额叶重整合三个环节的差异中，逼出“尺度钝化”这一命名的发生学来源；其次，通过与主体性折旧论的正面对质，确立尺度钝化不可被还原为“人力资本折旧的一个极端特例”的独立理论边界，并以一个两期决策的核心权衡模型为其提供自锁均衡的形式化分析；最后，将诊断推至发展经济学的宏观维度，指出 AI 辅助下的跳跃式发展可能使追赶者永久性地越过内生尺度代际再生产所必需的“历史痛感区间”。文章以“拮抗性摩擦”为支点，探讨在新的历史条件下重新安置尺度再生产装置的制度条件，并给出可观测代理指标与整个命题的证伪条件。

关键词：关键词：尺度钝化；掌尺回路；内生尺度；主体再生产；AI 经济学

一、引言：在决定的背面

这篇论文处理的是一记位移。这记位移至今尚未在经济学中被正式命名，但它正在沉默地改变“经济主体”这个概念的实质质地。

让我们从一个反直觉的经验观察切入。近年来，在金融投资、创业决策、公共政策制定等高度不确定的领域中，出现了这样一种现象：一批受过精英教育、能熟练调用最强 AI 决策辅助工具的年轻从业者，在面对那些超出所有历史回测与预研范围的极端情境时——一场从未有过的市场断裂，一次三线并发的公共卫生与就业危机——他们中相当一部分人做出判断的速度、质地与事后二阶反思的深度，不仅没有随工具升级而提升，反而低于那些在早年亲手经历过严重失误并独自承担后果的前辈从业者。

这并非一个关于“AI 让人变笨”的通俗叙事。这些年轻从业者的知识储备、信息优势和分析能力，在所有可观测的维度上都优于前代人。他们缺少的不是知识，不是工具，甚至不是“判断力”这个空洞的字眼。他们缺少的，是一种更底层的、必须用自己的全部生命史去亲手重制的东西——那杆借以称量“什么值得冒险、什么必须守住、什么可以放弃”的内生尺度。

这杆尺度不是在课堂上被教授的，不是在 AI 输出的光滑报告中被阅读的，不是可以从任何人的经验中提取为可迁移知识的。它唯一的生成原料，是主体在亲自面对一个无法委托给任何他人、任何机器、任何模型的选项时，在不知结局的情况下用自己的肉身和全部未来下注，在后果真实地打回自身时将那一痛烙进脊髓层面的价值标定过程。本文称这一完整过程为掌尺回路，称那个在此回路中被生产出来的、不可转让的尺度为内生尺度。

当前 AI 对经济系统的深度嵌入，正在引发一记比“人被替代”更隐蔽的根基位移：不是人在经济决策中的角色被边缘化，而是掌尺回路在人的身体内部被悬停——人仍然在“做决定”，但那杆秤的再生过程已经不再被激活。本文称这种回路悬停为尺度钝化。

这一诊断试图逼出经济学尚未系统追问过的一个问题：当整个经济系统的决策效率飙到历史最高时，还有多少人在亲手掌尺？当一代人的内生尺度是在一个被算法预先消化过的光滑环境中生长出来的，他们还能在多大程度上，在那些算法失效的时刻，用自己的身体去重新称量这个世界？

本文的论证分六步展开。第一步，通过对一个具体的资本配置决策在两种模式下的神经级因果机制进行对比推演，让“尺度钝化”从决策现场的肌体过程中被逼出来，而不是预先定义好再贴上去。第二步，在此发生学追踪的基础上，与既有进路——工具论、成本论、主体性折旧论——进行正面检讨，指出它们虽层层深入，却共同止步于“尺度是给定的”这一未被触及的先验。其中与主体性折旧论的对质，是确立尺度钝化不可还原性的关键战场。第三步，建立尺度钝化与人力资本折旧、干中学、默会知识三个邻近概念的不可还原边界。第四步，给出一个两期决策的核心权衡模型，为掌尺回路的钝化陷阱提供自锁均衡的形式化分析，并给出陷阱不出现的参数条件集合。第五步，将诊断推至发展经济学的宏观维度，提出“历史痛感区间跳跃”假说。第六步，探讨拮抗性制度设计的理论基础，并给出命题的证伪条件与可观测代理指标。

二、掌尺回路的接通与悬停：从决策现场逼出命名

一个概念如果不能在真实的决策现场找到其发生的机制，它就是一个漂亮的修辞，而非一个可供分析的理论命名。本章的工作不是举例说明一个已知的结论——相反，“尺度钝化”这个命名，将从两种决策模式下因果机制的逐层差异中被逼出来。

2.1 亲自掌尺：躯体标记装置的启动、写入与调用

考虑一个标准的资本配置决策：主体面对一个高度不确定的初创标的，必须决定是否投入一笔对其而言数额可观的流动资金。决策环境的信息是严重不完整的——历史回测数据平庸，管理团队的过往记录稀缺，市场对该赛道的态度两极分化。没有任何标准化模型能在这一情境下给出确定性的答案。

在亲自掌尺的模式下，决策过程沿着三个在时间上继起、在神经机制上接力的阶段展开。

阶段一：启动期——本体论摩擦的神经编码。主体在做出最终决定之前，经历了一段无可跳过的“悬而未决”的心理物理状态。这一状态在神经层面的对应事件是明晰的：前额叶背外侧皮层在面对多个相互冲突的选项时被高强度激活，前扣带回的错误监测系统等多个可能的风险后果之间持续发出冲突信号，岛叶从内脏状态——胃部持续紧缩、心率在特定念头闪过时轻微加速——中读取并整合为一种弥散但精确的“这个不太对”或“这个方向虽然算不出但必须试”的躯体感觉。这组神经事件的总和，构成了本文所称的本体论摩擦——那种“没有任何人能替我承担这个决定的后果”的紧张，在身体内部被实时编码为可被此后调用的躯体状态标记。

阶段二：消化期——躯体标记的写入。决策做出后，后果真实降临。如果后果是正面的，杏仁核与腹侧纹状体的奖赏回路会在结果确认的时刻释放一波高显著度信号，将该决策的情境-躯体状态绑定编码为“可复现的正向标记”。更关键的是负面后果。一笔实质性的亏损——不仅是一个负数，而是必须延迟购房首付、必须在深夜向家人解释为什么账户余额少了一大截——会在杏仁核高度激活的状态下，将当前情境标记为“高威胁事件”。海马在此期间对决策情境、内脏状态、时间节点的同步编码，形成一组可在此后任何类似情境中被自动提取的记忆痕迹。前额叶腹内侧皮层在事后反当中，将这次亏损从一桩孤立事件整合进主体的长程价值排序框架——它从此不再是“一笔钱亏了”，而是“我的胃从此知道，有些看起来对的事，底下可能有我上次没看清的裂缝”。这一过程，就是躯体标记的写入：它不是在认知层记下一条教训，而是在神经层焊上一道刻度。

阶段三：调用期——躯体标记的无意识告警。几周或几个月后，主体面对一个结构类似的新决策。在他尚未完成任何形式的认知评估之前，岛叶已经自动提取了上一轮写入的躯体标记，并将之映射为当下情境中的身体微反应——胃部轻微收紧、对某个选项产生了一种没有语言可以准确命名的抗拒。这就是达马西奥（Damasio, 1994）所说的“在推理之前就已经标记了某些选项”的过程。主体在此刻做出的决定，不是纯粹认知推理的结果，而是他的认知分析与他的躯体标记告警系统协同结算的产物。他从上一轮失败中得到的东西，不是一个写在备忘录里可以检索的要点，而是一个长在他身体里的、无需调取就能自动激活的警报线。

这一整串从启动期的本体论摩擦、到消化期的躯体标记写入、再到调用期的无意识告警的完整弧线，就是掌尺回路的一次完整接通。它的产物不是一项可被写成操作手册的技能，而是一道不可逆地标定

了主体自身价值排序的内生刻度。每一个亲身经历过一次完整掌尺回路的人，他在那个决策之后已经不是同一个决策者——不是他的知识增加了，而是他用以称量未来所有选项的那杆秤，从此多了一道刻痕。

2.2 外包给 AI：躯体标记装置的重重静默

现在考虑同一个决策在深度外包给 AI 辅助系统的模式下展开。主体将标的资料上传，指令模型“分析该项目的风险收益比并给出建议”。模型返回一份两千字的报告：行业对标、竞品格局、团队量化评估、VaR 值回溯、推荐投资金额与最晚退出时点。主体阅读报告，确认信息的逻辑自洽性，点击出款。

从神经经济学的角度审视这一过程，躯体标记装置在三个阶段上均未被唤醒。

启动期的绕过。从主体阅读 AI 报告到点击确认，决策发生的部位是他的视觉皮层、语言理解区和前额叶的认知评估模块。他没有经历那段“在两个无法用数据充分比较的选项之间悬而未决”的心理物理状态——因为模型已经替他完成了选项的排序与最优解的标定，呈现在他面前的是结论而非困境。岛叶没有被激活去读取内脏信号，因为从主体的身体内部来看，并没有一个“我这具身体正在进入一场没有退路的赌局”的紧急状态需要被读取。启动期的本体论摩擦——那个对回路唤醒最关键的初始条件——被绕过了。

消化期的空转。如果后果是负面的，主体当然会感到失望、挫败、甚至愤怒。但这些情绪的神经编码与亲自掌尺模式下的躯体标记写入有质的区别。区别在于：当主体清晰地知道“这个决定是AI做的”——即使在形式上他在报告上签了字——杏仁核在后果确认时的高激活不会将当前情境标记为“我必须为此负全责的、改变我未来行为准则的事件”，而是将之标记为“一次需要在下回调整AI参数的事故”。海马不会将本次亏损的内脏状态与决策情境做深度绑定，因为从他的身体视角来看，他从未真正“进入”过那个决策——他是在决策已经完成后，被通知了一个结果。前额叶腹内侧的反刍整合同样被削弱，因为“我回看我的判断哪里出了问题”和“我回看AI给我推荐的东西哪里有问题”是两种完全不同的认知操作：前者是方向级质询，后者是操作级调试。消化期没有躯体标记被写入，是因为在那整段时间里，那场亏损从未被他的身体真正承认为“我的”。

调用期的失能。下一次类似决策出现时，主体的岛叶没有可以被激活的库存——上一轮的亏损没有在他的胃、他的睡眠、他凌晨四点的惊醒里留下可以此刻被自动提取的躯体记忆痕迹。他依然可以阅读AI的新报告，可以以比上一回更谨慎的态度审视其中的参数设定，但那个在亲自掌尺者体内自动响起的“等一下，这个方向的感觉跟上次我栽的那个坑很像”的无意识告警，在他这里是沉默的。他不是不谨慎。他是没有那根可以自动弹出来的、长在身体里的警报线。

2.3 悬停的定义与命名的发生

至此，“尺度钝化”可以从上述因果机制对比的内部被正式逼出。

这个命名不是在“观测到某些现象”之后被赋予的一个标签。它是在既有概念——人力资本折旧、干中学、默会知识、主体性折旧——在以下三个点上依次被打回来之后，被迫给出的新名字。

第一击：既有概念无法命名那个被绕过的东西。人力资本折旧处理的是技能的贬值，干中学处理的是经验对生产率的贡献，默会知识处理的是不可言传的技能习得。它们都在各自的层面上有效，但当面对“主体在外包决策时，被绕过的不是一个技能或一类知识，而是一个必须由启动、写入、调用三阶段构成的完整神经回路”时，这些概念没有字可以用。

第二击：主体性折旧的时态叙事无法容纳“从未被写入”。这是最逼近的那个概念。主体性折旧论已经推进到了“亲自犯错是再生产决策主体性的生产活动”这一层，它的深刻之处在于将“犯错”从消耗重新定义为了投资。但它的整个叙事框架——无论拓展到何种程度——都依赖于一个“折旧”的隐喻：有一个曾在此处的东西，现在磨损了，需要修复或重新投资。这个隐喻在面对“某一代人在其躯体标记写入最活跃的窗口期里，从未被逼着亲身进入过一场没有退路的赌局”这一状况时，在语法上就失效了。折旧预设存有曾在此处。钝化面对的是：此处的存有从未被发生出来过。这不是一种更严重的折旧，而是折旧这个概念无法进入的另一类问题。

第三击：修复方案的失效逼出新命名。主体性折旧论所依赖的政策路径——提供更多犯错机会、恢复实践空间、把决策权交还给个体——预设了一个前提：一旦机会被恢复，那个已经磨损的机制可以被重新激活。但第二章的机制推演表明，掌尺回路的沉默不是一次性的关机，而是一个自锁过程：钝化本身在推高重新启动的成本。对于一个从未在胃里焊进过刻度的主体，给他一个犯错的机会，不等于他就知道如何用自己的身体去犯这个错。那个启动回路所必需的“我亲自承担不可逆后果”的紧张，对他而言不是一种熟悉的旧感受被重新唤起，而是一种从未被写入过的神经经验——它的首次启动，需要穿过一堵他此前一直被技术舒适地保护着、从未被逼着正面撞过的墙。修复方案的失效证明：这不是一个旧的善被破坏了需要修复的问题，而是在新的历史条件下，一个生产主体本身的装置，正在被系统性地搁置。这需要一个新的命名。

定义：尺度钝化指一个经济主体的掌尺回路——其在内脏与神经层面亲身感受、消化并重制价值刻度的完整过程——因系统性缺乏“亲自承担不可逆后果”这一启动条件，而在启动期被绕过、在消化期未写入、在调用期无法激活的状态。它不是尺丢了，也不是尺磨损了；它是尺的生产线——那个靠本体论摩擦驱动、靠躯体标记作原料、靠每代人的肉身从头重建一次的工厂——被静默了。

三、与既有进路的分层检讨

在尺度钝化从前述因果机制追踪中正式逼出之后，它与既有进路之间的理论距离变得可以精确丈量。本节由远及近，依次检讨工具论、成本论与主体性折旧论，指出它们虽层层深入，却共同止步于一个未被触及的先验：尺度是给定的。

3.1 工具论与成本论：尺在谁手不是问题

将 AI 视为“通用目的技术”的工具论传统，基于一个它从未审视过的暗默前提：那个在新技术面前做决定、衡量选项、分配资源的主体本身，是给定的。任务模型可以穷举被替代的岗位、计算替代弹性、预测新创造的工作类别，但它在数学上要求主体的决策函数——那套在选项间排序的偏好结构和判断能力——在模型中作为外生给定或按固定规律演化的变量。技术在变，菜单在变，但看菜单的那个人仍然站在技术变革的洪流之外。尺在谁手里、尺是怎么被造出来和重制出来的，工具论不问，因为它的分析起点就是那把尺已经稳稳地握在主体手上。

成本论比工具论进了一步。它注意到了 AI 在极致压缩交易成本的同时，也消灭了一种被经济学长期默认为“生产条件”的东西——人通过亲自犯错、亲自消化后果、亲自在疼痛中修正而长出来的判断力。它将“犯错”标记为一种被忽视的隐性投资，并呼吁经济学将其纳入损益表的核算框架。但成本论的整个分析仍在一个未被触动的核算单位内运行：它试图为旧的决策主体补上一笔被漏掉的开支，却没有意识到，那个主体的构成过程本身正在因为“不再亲自犯错”而从内部被改变。它在为旧单位找新账，而真正需要追问的是：核算单位本身是否需要被重新生产？

3.2 与主体性折旧论的正面对质

近年来的分析将诊断推到了更深一层。这类论述——可概括为“决策主体性在持续外包后发生磨损”——的核心洞见是：“亲自犯错”不是一种消耗，而是一种生产活动；它不是对已有主体性的支出，而是新主体性的生成。这一命题将分析焦点从技能的贬值移向了做决定者本身的质地的改变，是目前既有进路中最逼近尺度钝化诊断的那个位置。

但这里有一个必须被精确暴露的缝隙。

主体性折旧论所依赖的整个隐喻框架——“资产”“折旧”“磨损”“修复”——共享一个本体论预设：在某个初始时点，存在过一个相对完整、相对健全、相对活跃的决策主体性状态。这个预设使得该框架在解释一代人的主体性退化时有效（他们确实曾拥有过较高的内生尺度），但它在这一代人退出经济舞台之后，面对一个从未被逼着亲自进入过一场本体论痛感事件的年轻世代时，在语法上就失效了。你无法折损一个从未被写入过的东西。你无法修复一条从未被建造过的生产线。

这不是一个程度的问题，而是一个类的差异。主体性折旧处理的是：存有曾在此处，现在磨损了。尺度钝化处理的是：此处的存有从未被发生出来过，而那个使存有得以被发生的生产条件本身，正在被当前的制度与技术环境系统性地搁置。两者的差别，不是资产的存量为零和存量为正然后递减的差别——这种理解仍然在折旧的句法内部打转——而是“存有一亏缺”这一整套叙事无法进入“发生—未发生”这一套叙事的差别。

用一个类比可以更清晰地显影这种差别。人力资本折旧可以解释为什么一个五十岁的 COBOL 程序员在 Java 时代失去了谋生能力——他的技能曾经有效，现在贬值了。但人力资本折旧无法解释为什么一个从未被教过任何编程语言、从小在完全不提供这种教育的环境中长大的人，不会编程。后者不是“折旧”，是“从未形成”。将后者称为“折旧的一个特例”，在语言上可行，但在理论上混淆了两种完全不同的因果路径：折旧的因果链是“环境变化→存有贬值”，而从未形成的因果链是“条件缺失→存有未发生”。把“条件缺失”等同于“环境变化”的一个子类，就是把生产条件的撤销等同于市场的正常波动——这恰好是那笔被成本论和主体性折旧论共同漏掉的账。

尺度钝化命名的是一个更具体的“从未形成”：在二十岁到三十五岁这个躯体标记写入最活跃的窗口期里，一代人几乎没有被逼着进入过一场必须亲自在极度不确定中下注、并独自消化全部后果的本体论地带。他们的内生尺度不是磨损了——它从一开始就没有以躯体标记的形式被充分写入。他们所读过的所有案例、所有模型输出、所有关于风险的教科书解释，都储存在他们的前额叶和海马的可陈述记忆里，但没有储存在他们的胃里、他们的窦房结里、他们凌晨惊醒时那股说不清原因但知道“不能再这么下去了”的紧急感里。

3.3 跨界对话：内生偏好、默会知识与躯体标记

在进入形式化分析之前，必须与三个强大的邻近概念完成简要的不可还原边界勘定。

内生偏好。贝克尔和斯蒂格勒将表面变化的偏好看作不变基本欲望在不同价格与收入条件下的消费表现；比辛和维迪耶进一步将偏好形成内生于文化代际传递。内生尺度与内生偏好之间有一个范畴性的差异：偏好处理的客体是“如何排序给定选项”，而内生尺度处理的客体是“使排序能力本身得以被生成和重制的前提条件”。代际的偏好传递——下一代通过观察上一代的行为来更新自己的偏好——预设了传递的方向与对象。而内生尺度的产权属性恰恰在于其不可转让性：上一代人的胃曾经在亏损时紧缩过，不能直接把这道刻度的神经编码拷贝给下一代。下一代人必须亲自亏损，亲自让杏仁核和海马在那场亏损的高唤醒状态下写下属于自己的那组躯体标记。代际传递的不是尺度本身，而是——在最好的情况下——上一代人是否有意识地在制度中为下一代保留一个可以安全启动掌尺回路的摩擦空间。内生偏好解决的是“学什么”的问题，内生尺度解决的是“学本身是不是被开启了”的问题。两者不同层。

默会知识。波兰尼指出我们知道的远比我们能说出的多——骑手在高速入弯时调整重心的方式无法被写成手册，只能通过“寓居其中”的亲身实践来习得。掌尺回路在“不可形式化”这一点上与默会知识同构，但两者的产出具有范畴性的差异。默会知识的产出是一项可重复展现的技能：你学会骑车后，可以在任何一辆类似自行车上自动调用它，而不需要在每次上车时重新做出一场“为什么要骑车”的存在论决断。而掌尺回路的每一次完整接通，产出的是一个不可逆地标定了主体自身价值排序的、用于做出下一次唯一性决断的尺度。骑车的技能不会改变你是谁，而掌尺回路的每一次被后果打穿，都在改变那个做决定的人本身。默会知识生成的是“知如何”，掌尺回路生成的是“成为谁”。

躯体标记假说。达马西奥的躯体标记假说是本文所借用的最重要的理论地基，但本文在地基之上加建了两层。第一层：躯体标记不仅是决策的辅助工具，它本身是主体在漫长生命史中亲手锻打价值刻度的生产装置。那些被胃记住的疼痛、被睡眠中断记下的惊惧、被对整个家庭命运的担当所绷紧的神经，不是在“帮助”主体排序选项——它们就是那个排序得以被重制出来的唯一工厂。第二层：这个

工厂有一个被经济学严重忽视的脆弱性。它的活跃需要一种特殊的工作条件——主体必须反复进入“这个后果是我的，没有任何第三方可以替我吃下”的现实，才能把每一轮经验转化成可以被躯体存储并在此后自动激活的标记。躯体标记假说描述了装置在运作时的机制，本文追问的是装置的启用与维持条件本身——当这些条件被系统性绕过时，装置不是“出错”了，而是根本不在线。

四、不可还原边界：与三个概念划清界线

在完成与主体性折旧论的正面对质之后，尺度钝化作为一个独立理论对象的边界已经显露出了主要轮廓。本节进一步与三个更易混淆的既有概念划界，完成其理论边界的确立。

4.1 与人力资本折旧

人力资本折旧发生的位置是技能储备与市场需求之间的匹配程度。一个曾经的 COBOL 程序员在 Java 时代失去了谋生能力——他的技能存量贬值了。尺度钝化发生的位置比技能层深一层：它在掌尺回路本身。

一个经历了严重尺度钝化的企业高管，其人力资本几乎不会贬值。他的行业知识、分析框架、在 AI 辅助下甚至可能比此前更显厚实。但他作为“一眼看去就知道这个收购案哪里不对”的最终拍板人——那“一眼”背后所需要调用的全部躯体标记——已经因为长期没有被真正的后果打穿过而静默了。他的人力资本在增值，他的内生尺度在钝化。两者不仅可以无关，甚至可能背道而驰。

4.2 与干中学

阿罗的干中学模型中，学习是生产过程的副产品，其自变量是累积产出。但这个函数在描述累积产出对单位成本的影响时，完全不触及那个学习主体本身是否仍在“我亲自为后果承责”的神经级确信中被每一批产出所激活。

一条建立在 AI 辅助之上的学习曲线可以极陡——主体可以在极短时间内掌握大量可陈述知识和分析框架。但升得越快的那个人，其掌尺回路反而越可能没被真正启动过。不是因为他不努力，而是因为他学习的大部分工序都不再需要那个“我亲自承受代价”的摩擦信号来提供躯体标记编码所必需的神经唤醒。尺度钝化不是“学习得不够”，是“学习时被用来编码躯体记忆的那个肉体基底被绕开了”。

4.3 与默会知识

默会知识与掌尺回路在“不可形式化”上处于同一光谱，但在产出的性质上分道扬镳。默会知识产出一项可重复的技能——一旦学会，自动调用，不需每次重新存在论决断。掌尺回路的每一次完整接通，产出的不是一个技能，而是一个不可逆地标定了主体自身价值排序的、用于做出下一次唯一性决断的尺度，且每一次都在改写那个做决定的人本身。前者是技能的肉身化，后者是主体性本身的肉身重铸。

五、钝化陷阱：一个两期决策的核心权衡模型

至此，“尺度钝化”已经从一个发生学命名推进到了一个具有明确因果机制的理论对象。本节给出一个概念性的形式化框架，展示掌尺回路作为一种独立于人力资本和经验积累的状态变量，如何在主体的终身效用最大化决策中被系统性地搁置，进而形成一个自锁式的“钝化陷阱”。模型的目的是提供可分析的结构，而非求取封闭解。

5.1 模型设定

考虑一个经济主体在两个时期 $t=1,2$ 中做决策。每期初，主体面对一个投资决策：他可以选择亲自决策 ($a_t=1$) 或委托 AI 决策 ($a_t=0$)。

状态变量。设 s_t 为掌尺存量——即主体借以在高度不确定情境中亲自称量选项分量的内生尺度之敏锐度。它是对躯体标记编码密度与回路传导效率的简约表征。其运动规律为：

若 $a_t=1$ （亲自决策）： $s_{t+1}=s_t+\eta \cdot I$ ，其中 $\eta>0$ 是学习率参数， I 是当期结果——无论成败——所携带的“本体论信息量”。成败都可能贡献正的尺度增量，区别在于量级而非方向。

同时，亲自决策需要支付本体论摩擦成本 $C(s_t)$ 。当 s_t 较高时， $C(s_t)$ 较低——掌尺是一种熟习的身体活动，启动成本随回路活跃度上升而下降。当 s_t 较低（钝化较深）时， $C(s_t)$ 显著上升——因为重新启动那条已被阻尼的回路，需要穿过一堵“首次启动的墙”。

若 $a_t=0$ （委托 AI）：主体获得 AI 给出的事前确定等价效用 π_{AI} ，无需支付 $C(s_t)$ 。但 $s_{t+1}=s_t$ ——掌尺存量不发生任何增量更新，回路未被接通。若连续多期委托，可设 $s_{t+1}=\delta s_t$ ，其中 $0<\delta<1$ 代表休眠状态下回路传导效率的自然衰减。

需要说明的是，模型中的 $C(s_t)$ 仅捕捉了启动掌尺回路所需支付的那部分可被主观体验为“痛苦”的成本。在实际的神经级层面，本体论摩擦是一个更完整的回路概念：它包括启动期的神经冲突编码、消化期的躯体标记写入、以及调用期的无意识告警生成。 $C(s_t)$ 是它的一个可观测代理，而非它的全貌。

效用函数。每期效用不仅取决于货币收益 π ，还取决于决策本身的状态： $U_t=u(\pi_t)-1_{\{a_t=1\}} \cdot C(s_t)$ ，其中 $u(\cdot)$ 是标准的递增凹函数。终身效用为 $V=U_1+\beta \cdot U_2$ ，其中 $\beta \in (0,1)$ 为跨期贴现因子。

核心权衡。主体在每期必须权衡：选择 $a_t=0$ 可以在当期获得一个不差且光滑的货币收益，完全免去 $C(s_t)$ 的痛苦，但代价是放弃 s 的提升；选择 $a_t=1$ 可以在长期提升自己的内生尺度存量，使未来的亲自决策成本降低且质量提升，但必须在当期承担 $C(s_t)$ 这一真实的、在身体上被体验到的摩擦成本。

5.2 钝化陷阱的自锁均衡

假设主体从一个较低的 s_1 出发——比如一个在整个青年时期几乎从未被逼着独自承受不可逆后果的主体。他在第一期面临的 $C(s_1)$ 很高，因为他的身体对“亲自掌尺”的体验是完全陌生的。即使第二期的 s_2 提升可以带来跨期收益，该主体也极有可能在第一期选择 $a_1=0$ ——因为当期的摩擦对于他而言过于痛苦。

但一旦 $a_1=0$ 被选择， s_2 就停留在 s_1 的水平（甚至更低）。那么第二期他面临的将是同等甚至更高的 $C(s_2)$ ，于是他再一次选择 $a_2=0$ 。

这个锁死的关键在于：钝化本身在推高重新掌尺的成本。钝化越深，回去掌尺的价格就越高，而再回去掌尺的动力就越小。这并非一个偏好的“短视”可以解释的问题——从主体的视角出发，他每一期的选择在当期痛苦水平下都是个体理性的。但这一串个体理性选择的总结果，是一个在本体论层面自我加强的恶化循环：钝化的人在理性地逃避那件唯一能阻止他进一步钝化的事。

反之，一个 s_1 本来就较高的主体，在每期面对一个较低的 $C(s_t)$ ，因此选择 $a_t=1$ 的门槛较低， s 继续提升， $C(s)$ 进一步降低——形成正向螺旋。

5.3 陷阱不出现的参数条件

钝化陷阱是“在 AI 时代几乎必然发生”的普遍趋势，还是“仅当某些参数条件满足时才会发生”的特定均衡？模型的设定决定了它是后者。让钝化陷阱不出现的条件集合如下。

条件一：极高的跨期贴现因子（ β 接近 1）。如果主体极度远视，把“早日成为能掌尺的人”的终身效用折现值看得极重，那么即使当期的 $C(s_1)$ 很高，他也可能咬住牙在第一期亲自掌尺。这个条件的现实对应物是：一个社会能否让年轻从业者在职业早期就清晰地看到“不掌尺的终身后果”——例如，让他们观察第一代把全部决策权外包给模型的人在中晚年展现出的本体论质地的差异——从而改变其有效贴现因子。

条件二：低初始门槛 $C(s_0)$ 。如果 s_0 不是从零出发，而是在正式进入高风险决策之前，就通过一系列无害的、后果可控的“微掌尺”训练（如无 AI 辅助的小额投资、独立管理一个小预算项目并独力承担结果）建立了基本的回路传导效率，那么 $C(s_1)$ 从一开始就不会高到足以锁死整个链条。这是教育政策最直接的着力点。

条件三：委托 AI 决策本身带有显著的当期负效用。如果“把决定权交出去”这个动作本身——不是因为结果不好，而是因为“交出去”这个动作在本体论层面让主体感到一种无法忽视的自我背叛——带有足够大的负效用，那么 $a_t=0$ 的有效净收益就不再是光滑的 π_{AI} ，而是 π_{AI} 减去一笔内生的“委托疼痛”。这笔疼痛的来源不是外在的制度惩罚，而是主体在长期被训练（或文化性地被塑造）为“掌尺是人之为核心的活动”之后，每一次外包都在身体内部触发的一种深刻的违和。这个条件的现实对应物居于文化层面：一个社会能不能在对技术进行拥抱的同时，并行地维系一套“亲手掌尺是高贵的”文化编码，使得委托本身在主体内部产生一道无法被光滑过去的阻力。

这三个条件各自对应于政策、教育与文化三个维度上的制度着力点。后文的拮抗性设计，正是从这三个参数条件出发的推展。

六、跃过痛感区：追赶中的代际钝化风险

模型层面的钝化陷阱是个体层次的，但这一机制一旦乘以“代际”与“国家追赶战略”两个乘数，其宏观含义就变得不容忽视。本节提出一个此前未被系统追问的假说：后发国家在 AI 辅助下实现的跳跃式发展，可能使一代人永久性地越过内生尺度代际再生产所必需的“历史痛感区间”。

发展经济学对后发优势的传统叙事是清晰的：追赶者可以省去先行者走过的弯路，压缩技术学习和制度试错的成本，在一个较短的窗口期里完成先行者花费数十年乃至上百年来完成的工业化与制度现代化过程。AI 技术把这种“压缩”推至了历史极限：一个非洲或南亚的年轻人，现在可以直接在手机上调用全球最先进的大语言模型来撰写商业计划书、分析市场、规避合规风险——他不需要经历上世纪九十年代硅谷创业者那种从车库、失败产品、被银行拒绝、被合伙人背叛中一寸一寸用肉身摸出商业嗅觉的过程。

但这里有一个被“压缩成本”叙事遮蔽了的结构性危险。发达国家在其漫长的工业化与技术革命过程中，虽然支付了巨大的摩擦成本，但这种“慢”无意中为每一代经济主体的掌尺回路提供了一种天然的激活节律。市场崩溃是渐进发生的，危机过后有数年恢复期，企业家在破产后有社会缓冲去重新站起来，司法体系对商业欺诈的惩罚是逐渐累积的。这些事件不是“好事”，但它们以一种无可跳过的、绑在实在经济周期上的节奏，为每一代经济主体提供了不可替代的躯体记忆编码原料。那些胃紧缩过、睡眠中断过、在对员工宣布解散时嗓音发抖过的身体，一笔一画地，在几十年的职业生涯中刻出了对“什么是危险的、什么是账面上的亏、什么是致命的”这三者差别的内生尺度。

AI 辅助追赶的巨大诱惑在于，它可以同时压缩两类东西：一类是纯粹的“弯路”——技术试错与制度摸索的冗余成本，这类压缩是后发优势的实质；另一类是夹在弯路中间的那些“本体论痛感事件”——企业家在极度不确定中独自扛起全部责任、被失败打穿后重新校准自身价值排序的那些事件。这类“压缩”实际上不是节约，而是对一种无法被外部输入的生产要素的消除。当追赶型经济体用政策全力推进 AI 在各个经济决策层中无缝嵌入时，它极易在无意中创造出这样一代年轻精英：他们在二十岁到三十五岁这个躯体标记写入最活跃的窗口期里，几乎没有被逼迫着进入过任何一个必须亲自在极度不确定中下注并独自消化全部后果的本体论地带。他们可以条理清晰地表述最优风险管理策略，但他们的胃从未在真正亏损时紧缩过，他们的睡眠从未被那个因为自己判断失误而发不出薪水的员工的面孔打断过。他们不是没有知识，他们是身体里从未被烙进过那些知识的切身刻度。

这就构成了一个深层的代际再生产断裂：上一代人的内生尺度是从匮乏与痛感里自己长出来的，但那个刻度的产权恰恰在于它不可被移交。它不能被写进教科书、不能被打包进 AI 训练数据、不能通过任何代际间传递机制从上一代完整移交到下一代手中。下一代在 AI 的平滑环境中长成，他们继承的是一整套可以即刻调用的世界级分析工具，但他们没有继承——也无法继承——那杆必须亲自去锻打才存在的秤。

在此框架下，二十世纪后半叶多个经历过极速压缩式工业化的东亚经济体所呈现出的第一代创业者与第二代、第三代继承人之间在面临系统性转型时判断力质地的差异，可以从“历史痛感区间是否被越过了”这一角度被重新审视。关于“二世症候群”的传统叙事往往聚焦于“接班人的能力不如创始人”，而本文提供的假说将问题推到另一个层面：不是第二代的“能力”更差，而是他们成长期间所

处的制度与经济环境，从未向他们提供过他们的父辈所曾经历的那种级别的本体论痛感事件。他们不是在能力上更弱——他们是在躯体记忆的密度上，远远薄于他们的父辈。而当一个经济体在产业升级或跨周期危机的关口需要一个能在所有模型失效时挺身做出突破性判断的人时，那个人的“躯体标记储量”可能比其“认知能力”更关乎成败。

如果这个假说成立，那么发展经济学在其核心的追赶逻辑中引入一个新的逆向筛查指标就变得必要：不是问“AI 渗透率是否够高”或“技术采纳速度是否够快”，而是问——这一代年轻创业者在三十岁之前，是否有过至少一次有案可查的亲自决策并承担了超过半数的本金亏损？如果这一比率在 AI 辅助政策全面铺开之后出现长时期的系统性下降，那么极有可能发生的情形不是“创业成功率上升了”，而是亲自承受痛感的机会正在从一代人身上被政策性地消除，那一代人的掌尺回路，正在整体性进入钝化通道。

七、拮抗性设计：在矛盾中安置再生产装置

至此，本文的诊断已经完成。但它若仅止于此，就是一份警告，而非一个可以展开讨论的责任框架。必须追问：在人机深度交互不可逆的历史条件下，是否存在制度设计的可能，把那些被 AI 的效率逻辑判定为“冗余摩擦”而试图消除、却又恰恰是掌尺回路唯一激活信号源的本体论摩擦，做成不可被卸载的保留装置？

这需要一种与当前技术应用主流逻辑拮抗的实践伦理——不是反对 AI，而是在 AI 全面接管了工具性决策摩擦之后，反过来刻意保留，甚至是有意创造那些无法被外包的“直接接触面摩擦”：人独自面对不完全信息、承担不可逆后果、必须亲手在几个无法被数据充分比较的选项之间称量分量的那种本体论紧张。

从第五章模型给出的三个参数条件出发，拮抗性设计可以在三个层面上着力：降低初始启动门槛（ $C(s_0)$ ），让掌尺回路在不致命的低压环境中被首次接通；拉高跨期贴现因子（ β ），让年轻从业者尽早看见钝化的终身后果；和注入委托疼痛，让“把决定交出去”这个动作本身带有不可被忽视的负效用。

教育层面：低压微掌尺训练。在大学教育的前半段，设置必须完成的“无 AI 独立探究模块”，给学生一个无法通过查询模型就找到最优解的开放问题——例如“本地摊贩管理条例的实质目标与隐性代价是什么”——要求他们在一个学期内自行定义核心问题、独立完成实地观察与原始数据采集、从互相冲突的证据中建立论证、最后在公开答辩中为自己的立场辩护。这个过程中不可避免的大量微观错误——选错观测站点、设计出无法在现实中执行的问卷、在窗口期临近时才发现自己最初的方向是歪的——不是在产出那篇论文本身，而是在产出另一件东西：一个将来无论他读到多少质量优良的 AI 报告，都能身体级地怀疑其中“光滑得可疑”的部分的内生警报线。这一设计的参数学含义是降低 $C(s_0)$ ：在后果可控的低压环境中，让主体提前完成第一次掌尺回路的接通，从而在高风险决策到来之前，回路已经具备了基本的传导效率。

职业制度层面：硬性独立判断节点。一个已经在某些金融机构的自发实践中被观察到的做法，可以作为一个拮抗性设计的微型前例。其操作形式是：在投资委员会对一个重大标的做出正式决策之前，硬性规定每位必须签字的决策成员先在隔绝 AI 辅助报告的条件下，手写出一份不短于八百字的独立判断备忘录，并亲笔签名。这份备忘录不是用于评估该标的本身，而是用于在事后将该决策成员的亲身判断与 AI 的判断做反向比照：如果事后发现某一项重大风险在 AI 报告中已被明确提示、但该成员在独立备忘录中完全没有触及，且决策最终因该风险而遭受实质损失，那么该成员所享有的职业责任豁免将受限。这个设计的逻辑不是“把决策变慢”，而是把“我必须为我自己写在纸上的话担责”这种本体论摩擦重新植入一个已经被 AI 光滑输出覆盖掉所有紧张感的流程里。它的参数学含义是提高委托的当期负效用：让 $a_{t=0}$ 这个动作——把判断权交出去——本身在一个制度性的独立判断节点面前，产生一种内生的、不可被忽视的“委托疼痛”。

国家层面：宏观逆向监测指标。正在大规模将 AI 纳入产业与发展政策的发展中经济体，可以参照一个不设目标值、只观察下降趋势是否超出预设底线的逆向指标——“年轻创业者三十岁前亲自决策并承担过半数本金亏损的比例”。如果这一比例在政策推动 AI 创业辅助系统后出现持续性、不可逆的

下降，这可以被视为一个比 GDP 增速下滑更需要警惕的宏观信号。它意味着不是“创业环境变好了”，而是痛感——那个无法被定价、无法被外包、无法被代际传递、却又是经济主体最为稀缺的再生产原料——正在从这一代人的身体里被系统地抽走。

这些拮抗性设计不是对旧时代决策方式的浪漫化回归。它们要做的不是“守护一个被 AI 威胁的旧善”，而是承认一个不可逆的历史条件——AI 已经永驻经济系统的决策底层——之后，在它的逻辑内部，重新安置那些生产经济主体自身所必需的摩擦。这不是回归。这是从矛盾里逼出新的生产条件。

八、结论与证伪

本文从头至尾追着一个质问：当经济系统把判断的准备工序外包得越来越干净时，那个在做决定之前必须亲自重制出借以判断轻重之尺的掌尺过程，是否正在从这一代经济人的身体里被切断？

结论是，它正在被切断。这个断口的名字是尺度钝化——它与人力资本折旧无关（尺的钝化可以与人力资本的增值同步发生），与干中学不同源（学习曲线可以陡峭攀升而尺从未被真正调用过），与默会知识不同质（后者生成可重复的技能，前者生成不可逆的主体性重铸），与主体性折旧不同位（后者预设存有曾在此处，钝化面对的是存有从未被发生出来过）。

在 AI 时代，一种健全的经济体——不是 AI 渗透率最高的那一个，而是那个在技术把绝大部分工具性摩擦都压缩殆尽之后，仍然有意在制度中为“掌尺回路”安置微型激活窗口的经济体——将不是在守护旧的人，而是在持续地、以拮抗的姿态，在人与技术的深度纠缠内部，生产着下一代还能亲手掌尺的人。这不是回归到某种人的本质，而是直面一个不可逆的矛盾之后，在它的结算中逼出新的人的生产条件。

本文的全部命题必须被置于可被驳倒的条件下才能获得体面的学术身份。

可观测代理指标。“尺度钝化”作为一个理论概念，在直接测量上存在困难——我们不能切开一个人的颅骨去测量其岛叶-腹内侧前额叶回路的传导效率。但它有三个可被独立观测的行为代理变量：

（1）分辨退化指标——在实验条件下，钝化者能否在多个 AI 输出的光滑选项中，对某个选项产生身体级的抗拒（如心率变异性变化、皮肤电导反应），即使该选项在认知层面无法被反驳；（2）决策委托阈值的协同性变化——通过跟踪调查，观测“这个决定太难了，还是让算法来吧”这一态度在实际决策行为中出现频率的变化是否与问题的客观复杂性脱钩；（3）二阶自省的方向层级——在对失败决策的事后反思中，钝化者的归因更倾向于“下次应该调整AI的输入参数”（操作层）还是“我最初进入这个决定的大方向是不是在根基处就是歪的”（方向层）。

证伪条件。如果在至少两个文化背景不同的发达经济体中，那些已将 AI 决策辅助全层嵌入核心决策环节并已连续实施超过一整代人（不少于二十五年）的组织与社会部门，在面对一次超出所有历史回溯数据范围、所有预研场景均未覆盖的链式极端决策危机时，不仅常规应对速度显著优于依赖非 AI 并行培养体系的对照组，而且危机中期的二阶自省表现也显著优于对照组——即独立于 AI 输出的阶段性简报，有相当比例的决策者能主动质询“我们最初被数据推入的这个大方向，是否在根基处就是歪的”，并在后期成功实质性扭转原定航向——那么，本文的核心诊断就被证伪了。因为那将意味着：掌尺回路的接通与维持，在一种彻底被模型外化的制度化协同中，可以完全以另一种形态——不需要每个主体各自经历的不可逆痛感事件——被集体性地保全甚至增强。也就是说，掌尺本身不再是不可被制度外化的终极因素，内生尺度也不再必须是“内生的”。

如果在可见的未来，这一证伪条件没有被触发，而本文的诊断所警示的那种代际钝化仍在以沉默的方式持续——那么，这篇论文所做的工作，就是在它尚可被追问的时刻，为经济学这盏以人为主语的灯，在技术变局的深夜里，标出了一个它不可轻率放下的追问：当所有模型都失效时，还有多少身体里仍燃着那根用自己的一生去亲手刻出来的尺？

参考文献

- Arrow, K. J. (1962). The Economic Implications of Learning by Doing. *The Review of Economic Studies*, 29(3), 155–173.
- Becker, G. S., & Stigler, G. J. (1977). De Gustibus Non Est Disputandum. *The American Economic Review*, 67(2), 76–90.
- Bisin, A., & Verdier, T. (2001). The Economics of Cultural Transmission and the Dynamics of Preferences. *Journal of Economic Theory*, 97(2), 298–319.
- Damasio, A. R. (1994). *Descartes' Error: Emotion, Reason, and the Human Brain*. New York: G. P. Putnam's Sons.
- Polanyi, M. (1966). *The Tacit Dimension*. New York: Doubleday.