

规范性自噬：AI时代知识主体性的系统性耗竭与假性外部生产

AI时代知识主体性的系统性耗竭与假性外部生产

陈晓艳 著 · 依王德生《SDE本体论》· 发表于2026年7月28日 · 创新智商待独立复核

摘要

人工智能的迅速普及正在改变知识生产的劳动结构，也使知识主体性问题重新变得尖锐。本文提出“规范性自噬”框架，用以解释一种具有条件性的制度悖论：当认知系统不断加密审核、方法训练和合規程序以维护判断质量时，这些维持性安排可能反过来削弱成员形成独立判断所需的探索经验。本文区分四个相互衔接的机制：互锁式认知空转、排斥形态由显性惩罚向担保虚化的转变、诊断知识在传播中的经验压缩，以及系统对可控批判位置的制度化生产。最后一种机制被命名为“假性外部生产”：组织表面上持续引入挑战，实际却可能预先限定挑战的议程、强度和后果，从而以批判仪式替代结构更新。本文将该框架与德里达的自体免疫、阿吉里斯的组织防御、默顿的目标置换以及组织学习理论进行比较，明确其重叠部分与可判定差异；同时提出同质多案例比较、过程追踪和教育准实验等检验路径。本文不把规范性本身视为病因，而把规范密度、任务新颖度、裁量空间与判断质量之间的特定组合关系视为待检验对象，由此把一项哲学诊断转化为具有经验风险的中层理论。

关键词：规范性自噬；担保虚化；认知劳动；假性外部生产；判断主体性；显影的自噬；知识生产

1 引言

1.1 问题的提出

2023年以来，以大语言模型为代表的生成式人工智能的迅速普及，深刻改变了知识生产的基础生态。学生用AI完成论文草稿，研究者用AI生成文献综述，程序员用AI辅助编写代码，政策分析师用AI提炼决策备忘录。在各个认知劳动领域，一个共同的现象正在发生：AI不再仅仅是“查找信息的工具”，它已经成为认知过程的深度参与者，甚至在很多情境下，成为认知过程的替代执行者。

一个值得警惕的张力由此出现。人工智能能够承担检索、归纳、改写和初步论证等任务，确实可能释放部分认知资源；但如果使用者在问题界定、证据权衡和结论担保等关键环节也持续依赖外部生成，其独立判断是否会受到影响，仍是一个尚待充分检验的问题。这里需要区分短期绩效与长期能力：任务完成得更快，并不自动意味着判断形成得更扎实；反过来，使用外部工具也不必然造成能力衰退，关键取决于哪些工序被外包、使用者是否保持核验责任，以及训练情境是否保留了必要的困难。

因此，本文不把问题简化为“AI导致技能退化”，也不把一切后果归因于技术。本文关心的是一个制度层面的条件性机制：当高度标准化的知识系统同时借助人工智能压缩探索、论证与修订过程

时，维护正确性的程序是否可能削弱产生新判断的经验基础。本文把这种潜在的自我侵蚀机制称为“规范性自噬”。该概念描述的是可被证伪的关系假设，而不是对所有规范、所有组织或所有AI使用方式的普遍断言。

1.2 核心判断

本文提出的核心判断可以凝缩为以下四个层次：

第一，在以可靠判断为核心目标的认知系统中，审核、方法训练、同行评议和合规检查具有不可替代的纠错功能，但它们并非只有收益。当程序密度持续提高、问题新颖度较高而个体裁量空间不断缩小时，这些安排可能把尚未成熟却值得探索的判断过早翻译为既有格式，从而挤压形成新概念所需的试错经验。本文讨论的不是规范与创新的简单对立，而是规范超过何种阈值后会改变判断生成方式的问题。

第二，这种排斥可能改变形态。较早阶段的制度通常通过拒稿、否定或边缘化等方式惩罚已经出现的异见；更成熟的制度则可能借助模板、指标和最佳实践，在判断形成之前就限定可接受的问题形态。后一种情形更难识别，因为成员未必感到被压制，甚至可能真诚地认为自己在自由判断。本文将这种从“惩罚异见”转向“削弱异见产生条件”的变化称为担保虚化：判断仍有署名和程序，却越来越少携带判断者亲自界定问题、承受不确定性并承担可错责任的过程。

第三，病理诊断本身也可能产生二阶效应。批判性思维课程、AI偏见清单和方法论警示能够减少重复错误，但当复杂的探索经验被压缩成可快速调用的规则时，后来者也可能只学会识别已命名的陷阱，而没有形成面对未知困难的能力。本文把这种风险称为“诊断知识的经验压缩”。它不是说知识传播必然削弱能力，而是提出一个边界条件：如果教学只传递结论和路线图，却取消独立尝试、失败和修订，诊断就可能在提高显性警觉的同时削弱迁移到新问题的判断能力。

第四，当真实批判越来越难以进入决策链，系统可能通过制度化的挑战者角色、外部咨询或创新活动维持“仍在开放”的形象。这些安排有时确实能够推动更新，但也可能成为一种可控的批判装置：议题由系统预先选择，批判止于不触动资源分配和责任结构的位置，参与者获得表达空间却缺少改变决策的通道。本文将后一种可观察的组织形态称为“假性外部生产”。判断它是否成立，不能依据组织是否举办批判活动，而应比较批判议程的独立性、意见进入决策的比例以及资源配置是否随后改变。

上述四个层次的判断，共同构成了一个具有时序形态（认知空转→担保虚化→显影自噬→假性外部生产）和清晰可证伪条件的理论框架，本文统称之为“规范性自噬”框架。

1.3 本文的结构

本文以下按标准学术论文结构展开。第二节在文献综述中，将本文的核心判断与五个高度相关的既有概念集群逐一对话：德里达的“自体免疫”、阿吉里斯的“组织防御性例行”、默顿的“目标置换”、医学免疫学中的“自身免疫”，以及莱维特和马奇的“能力陷阱”。本节将论证，上述每一个概念都在解释本文所关注的病理的某个侧面时不可或缺，却都因各自共享的一个未经追问的预设——系统拥有一个可被识别的“外部”——而无法完整覆盖本文所揭示的现象全貌。第三节构建理论框架，正式提出“规范性自噬”的三阶段机制，并显式界定跨学科理论资源调用的有效边界。第四节说

明本文所采用的概念分析方法论及研究设计。第五节为分析与讨论，逐层展开论证，并回应潜在的反驳。第六节总结全文，交代理论贡献、实践含义、研究局限，并基于可证伪检验方案提出未来研究的五个具体纲领。

2 文献综述

本文所关注的现象——规范性系统在维持自身秩序的过程中系统性地消耗其赖以生存的核心能力——并非全新的发现。在哲学、组织理论、社会学和免疫学等多个学科领域中，已有多个高度相关的概念集群分别捕捉了这一问题的不同侧面。本节的任务是将本文的“规范性自噬”框架与这些最近的学术邻居进行系统对话，在承认其贡献的同时，明确指出它们各自的盲区，从而锁定本文理论创新的不可还原性。

2.1 德里达的“自体免疫”：免疫性排斥的哲学结构

雅克·德里达晚期提出的“自体免疫”（autoimmunity）概念，为分析共同体如何在自我保护中损伤自身提供了重要哲学资源。德里达讨论民主、主权与宗教时强调：共同体为了保存自身而采取的防护行动，可能动摇其声称保护的开放性条件（Derrida, 2005）。这一思想与本文高度相关，但它首先是一种解构性的哲学结构，而不是关于组织变量及其因果效应的经验模型。本文因此只继承其“保护机制可能参与自我损伤”的问题意识，不把自体免疫视为一切规范系统的必然命运。

这一哲学洞见与本框架的亲性和是显而易见的。本框架的“规范性自噬”命名，其核心结构——“维持性排斥→自毁根基→诊断被卷入”——在形式上几乎是对德里达自体免疫结构的认知-规范领域的复写。然而，德里达的自体免疫分析存在两个关键局限，使之无法直接用于本文所关注的经验现象。

第一，德里达关注的是免疫逻辑内部不可彻底消除的悖论，而本文试图区分这种结构性张力在不同组织条件下是否会转化为可观察的能力衰退。二者的差异不能靠更强烈的隐喻来证明，而必须依赖比较研究：在规范密度相近的系统中，如果显性惩罚与程序性预塑对判断质量的长期影响没有差异，本文关于排斥形态的调节命题便得不到支持；如果程序性预塑在控制心理威胁后仍显著降低陌生问题上的独立判断质量，本文的形态区分才具有经验上的必要性。

第二，德里达的分析始终预设了一个边界清晰的“共同体”——自体免疫之所以可能，正是因为存在着一个可以被识别、被守护、被侵蚀的“自身”（self）。这一预设使他无法看见在本文所描述的担保虚化完成后的晚期阶段，当“自身”的边界已经弥散为纯粹规范性的场域时，免疫不再需要一个可识别的“抗原”（他者）来触发。系统不是在攻击“自身”，而是在掏空“能够做出自身/他者之分的判断能力本身”。在此阶段，自噬不再以“将自身组织识别为威胁”的经典免疫学模式运作，而是以一种弥漫性的、无特定对象的、在规范维持中自行发生的“基底掏空”模式运作。德里达的概念无法描述这一状态，因为他的免疫力始终需要一个边界来界定“何为自身”。

2.2 阿吉里斯的“组织防御性例行”：人际威胁回避的局限

在组织理论领域，克里斯·阿吉里斯关于“防御性例行”（defensive routines）的研究是理解“制度如何使聪明人无法学习”的经典贡献。阿吉里斯（Argyris, 1990）发现，组织成员为回避困

窘或人际威胁，会发展出一套高度熟练的防御性行为模式——例如，在会议上避免提出可能暴露自己无知的问题、用模糊语言包装尖锐的批评、或形成一种心照不宣的共识“某些问题不可讨论”。这些防御性例行系统地阻碍了组织学习，因为真正需要被审视的错误和假设被巧妙地绕过了。更为关键的是，阿吉里斯指出，对这些防御行为的讨论和诊断本身会激发更多的防御，使组织陷入一种不可纠正的陷阱。他创造“熟练的无能”（skilled incompetence）一词，来描述组织成员如何精于执行那些恰好阻止他们学习的行为。

这一框架与本文的亲性和性同样清晰。其“维持秩序（避免尴尬）→消耗有效判断→诊断被防御捕获”的结构，与本框架的“互锁式认知空转”和“显影的自噬”高度同构。然而，其关键局限在于，它将防御性例行的启动条件锚定于个体对人际威胁的心理感知和回避动机之上。在阿吉里斯的叙事中，组织成员之所以参与防御性例行，是因为他们感到困窘、感到被威胁、感到“如果说了真话会有不良后果”。这使得他的干预方案自然地指向了创造“心理安全”（psychological safety）和进行“双环学习”（double-loop learning）的对话。

但本文所提出的“担保虚化”机制，描述的恰恰是一种即便在心理安全极高、无人在回避困窘、个体完全自如地表达“合规判断”的环境中，仍然能发生的系统性判断力耗竭。在担保虚化状态下，个体不是不敢说话，而是他们所说出的任何话，都因为在日益精细化的流程中被预先格式化为“可供审核的合规选项”，而不再携带着那种未经消化的、粗糙的、充满认知张力的原初判断的质感。这种情况下，即便组织内部充满阿吉里斯式的“坦诚对话”，对话的内容本身已经是在规范的筛网中被过滤过的；人们真诚地交流着彼此真诚地认为有意义的判断，却未意识到这些判断之所以能“顺利通过”彼此的心理安全网，恰是因为它们的担保基底已经在更早的阶段被合规流程无声地挖空了。

若在一个组织内，成员普遍报告低人际威胁、高心理安全，但随合规流程精细化仍出现独立于回避动机的判断力显著下降，则阿吉里斯的概念无法解释，而本框架的担保虚化机制给出了替代解释。反之，如果所有判断力衰退均可还原为可测量的防御性推理与回避动机，那么阿吉里斯的框架足够，本文则是其多余的重述。这一可判定差异为未来的组织研究提供了一个清晰的经验检验方向。

2.3 默顿的“目标置换”：被偏离的“原初目标”何在

罗伯特·K·默顿在1940年代提出的“目标置换”（goal displacement）概念，是社会学对“规范维护如何走向自我挫败”的最早精确诊断之一。默顿（Merton, 1940）观察到，在官僚组织中，成员倾向于将严格遵守规则这一本来只是实现组织目标的手段的行为，逐渐当成目的本身。当这种情况发生时，组织虽然表面上秩序井然、合规度极高，但实际上已经偏离了其原初的功能目标，出现了“手段吞噬目的”的功能失调。

本框架的“互锁式认知空转”在形式上几乎是目标置换在认知维度的投射：系统为维持判断的规范性而建立的审核、分级、同行评议等制度，本应是为保护判断质量而服务的“手段”，但当这些手段在语言地形的空白区（即那些尚未被既有概念网络覆盖的新问题域）被严格执行时，它们实际上惩罚了那些试图为空白区造词的认知者，从而在维持规范秩序的同时，阻断了认知边界的突破。系统的维持手段（合规流程）吞噬了系统的原初目的（产生对真实问题的负责任的新判断）。

然而，默顿的目标置换框架存在一个在本文所关注的语境中变得致命的前提预设：它假定存在一个可被识别的“原初目标”，而“置换”之所以能被诊断，正是因为外界观察者可以将组织的当前行

为与那个原初目标进行比较，发现偏离。但当我们进入本文所描述的担保虚化完成后的晚期阶段，这一预设本身开始悬空。系统通过长年的维持性动作精细化，已经把“原初目标”的具体内容不断地重新解释、不断地适配当下流程的语言，以至于当成员试图追问“我们这个组织最初是为什么而存在”时，他们能够得到的回答，已经是一套在既有规范框架内被反复修剪过的、无害化的叙事。此时，置换不再是从A目标偏离到B手段，而是连A目标本身都已经无法被从那一套维持性语言中分离出来——系统维护的，仅仅是“仍有规则在运行”这一信号的持续广播。

比这更为关键的是，默顿的框架不包含“诊断行为本身被置换逻辑再捕获”的二阶反身性。在经典目标置换中，外部诊断者（如研究者、管理顾问）是可以识别出置换并提议纠正的，因为诊断者站在系统“原初目标”的规范性参照点上，能够比较“应该是什么”和“实际是什么”。然而，本框架的“显影的自噬”机制指出，在认知劳动的担保被虚化的条件下，诊断者本人也面临着同样的基底掏空——他所依赖的那些用于识别“原初目标”的分析工具和概念框架，本身可能已经是维持性规范的产品。他越是精于诊断，越可能在使用被规范磨平的概念工具时，将诊断成果压缩为一套可被传授的“反思方法论”，从而在更大范围复制了他所诊断的那种置换。

若在一个明显存在目标置换的组织中，外部诊断干预成功地恢复了目标导向，且未曾出现“以诊断之名强化更精细的空转规范”的后续效应，那么本框架的C机制（显影的自噬）即为伪，目标置换作为单一概念已经足够。反之，若诊断介入后，组织逐步以“我们已经反思过这一机制”为由，发展出一套更精致、更自我意识的合规体系，而这套体系在实质上继续抽空了成员的独立判断力，则说明存在一个超越目标置换的二阶机制在运作，默顿的框架无法充分描述这一后续，而本框架则直接将其纳入了理论的时序形态之中。

2.4 免疫学中的“自身免疫”：生物类比的有效边界

医学免疫学为本文提供的是受限制的结构类比，而不是社会机制的直接证明。自身免疫病涉及免疫耐受失调以及免疫反应对自身成分造成的组织损伤，其发生机制具有疾病类型差异，不能简化为单一的“把自身误认成外敌”（Davidson & Diamond, 2001）。本文借用的仅是一个形式关系：原本承担保护功能的机制，在特定条件下可能参与损伤其保护对象。社会系统没有抗体、抗原或生物学意义上的免疫耐受，因此后文所有“免疫”“自噬”表述都属于机制启发式概念，不能替代组织层面的因果证据。

本文的规范性自噬框架在纯形式结构上是对这一生物学机制的隐喻性复写：免疫系统（维持规范性秩序的审核、合规、评议流程）在识别和攻击“非我”（不符合既有范式的异见、无法被分类的新判断）时出现识别错误，将“自身”（系统赖以做出适应性判断的判断主体性这一核心能力）当作威胁加以攻击，导致系统的功能衰竭。这种形式同构是富有启发性的，但必须严格界定其作为类比的有效边界。

最关键的边界在于操作性定义的层次差异。在医学自身免疫中，无论是“自身组织”还是“免疫攻击”，都有明确的血清学、组织病理学指标作为操作性定义——自身抗体滴度、炎症因子水平、组织活检的病理改变——这些指标可以通过实验室检测被精确测量，其因果链条受制于生物学变量。而在规范性自噬中，“判断主体性”、“担保虚化”、“免疫记忆稀释”等核心变量的操作性定义基于社会认知层次——它们是制度性变量（合规流程的颗粒度、审核节点数、组织内“被允许提问”的空间）、认知变量（成员在面对模糊问题时的独立判断质量、对异质信号的敏感度）和行为变量（创新

尝试的频率、失败容忍度）的函数。其发生与消退不取决于个体的生理免疫状态，而取决于制度性规范的设计。

这一边界意味着，任何试图通过测量成员的生理自身免疫指标（如皮质醇水平、免疫功能）来直接“验证”规范性自噬的做法，都混淆了类比的层次。本文的证伪设计（见5.4节）因此将自身严格锚定在社会认知层次的变量操作化上。反过来，若某一社群或组织的判断力衰退，可完全由其成员的生理性自身免疫病流行率或慢性应激的生理指标所解释——即当控制了制度性变量后，判断力衰退随生理健康指标的变化而消失——则本框架为伪，医学模型已经足够。若判断力衰退随制度性变量（如流程审核节点数）而变化，且这一效应在统计上独立于个体生理健康指标，则本框架所特有的社会认知因果路径获得了支持。

2.5 能力陷阱：环境变迁是衰竭的必要条件吗

莱维特和马奇（Levitt & March, 1988）提出的“能力陷阱”（competency trap）概念，是组织学习领域对“成功导致失败”悖论的最经典表述之一。其核心逻辑是：组织因利用其现有能力获得成功而不断强化该能力——优化流程、积累经验、培训新人——在此过程中，逐渐减少对新能力的探索与尝试。由于在短期内，利用旧能力总是比探索新能力更有效率、更少风险，因此组织会在无意识中滑入一种对既有能力的过度依赖。当外部环境发生重大变迁（技术换代、市场重构、监管转向）时，因长期缺乏探索新能力而适应力萎缩的组织，将在新环境下遭遇失败。

这一框架共享了本理论“维持性优秀动作→掏空适应性判断力”的悖论结构，且两者都在解释一个“为什么会僵死在成功之中”的问题。然而，其关键差异在于，能力陷阱将衰竭效应的触发条件几乎完全锚定在环境变迁之上——旧能力之所以失效，是因为环境变了；如果环境一直不变，既有的优秀能力本可以继续成功下去。在能力陷阱的叙事中，坑在外部，组织是一脚踩进了环境变迁造成的陷阱中。

本文提出的规范性自噬框架则做出了一个更强的预测：即使外部环境高度稳定——技术轨迹清晰、市场需求无突变、监管政策长期一贯——系统内部的维持性动作的精细化过程本身，仍能内生性地消耗系统识别异质信号的能力。这是因为，维持性动作的精细化所磨掉的，不是“与旧环境适配的旧能力”，而是“识别环境中的异质信号本身所需要的认知感受力”。当一个组织的合规流程精细到令成员的每一次判断都必须经过多个节点的审核、每一次创新尝试都必须用已有规范的术语来论证其“不违反规范”时，即使环境中出现了真正的新信号，成员也已经失去了感知它的那种未经格式化的粗糙触觉。衰竭不是因为旧能力不适应新环境，而是维持旧能力的过程把“适应能力的生成基底”本身磨掉了。

这一差异具有极重要的干预意涵。如果能力陷阱是对的，那么组织的应对策略应当是关注外部环境变化，在环境变迁时及时“切换”能力。如果本框架是对的，那么即使环境看起来风平浪静，一个高度成熟的规范性系统也已经在内生性地丧失其应对下一次变迁的潜能——而当环境变迁真正发生时，该系统的崩溃将不是因为它“没有及时学会新能力”，而是因为它早已在平静中耗尽了能够“学会任何新能力”的那个学会者本身。

本文在5.4节设计的证伪方案，正是围绕这一关键差异展开：在控制外部环境高度稳定的条件下，检验维持性动作的精确度与内部判断力衰退之间是否存在独立的负向剂量-反应关系。若这一关系在环

境高度稳定时不成立，则能力陷阱的解释足够，本文的独立贡献被推翻。若这一关系在环境稳定时依然成立且呈剂量-反应模式，则意味着维持性动作本身具有一个独立于环境变迁的耗竭效应，能力陷阱无法覆盖这一经验区域，而本框架为此提供了理论解释。

2.6 既有研究的共同盲区：系统/外部的二元预设

经过上述与五个最近概念集群的逐一对话，一个共同的结构盲区浮现出来。德里达的自体免疫需要预设一个可识别的“自身”边界；阿吉里斯的防御性例行需要预设一个可感知的“人际威胁”作为回避动机的来源；默顿的目标置换需要预设一个外在于当前手段的“原初目标”作为偏离的参照；医学自体免疫需要预设一个可被血清学指标精确界定的“自身组织”；能力陷阱需要预设一个外部环境的变迁作为旧能力失效的触发条件。

在每一个理论中，系统的病理机制的运作，都依赖于一个相对于该机制而言是“外部的”给定物——外部的他者、外部的威胁、外部的目标、外部的组织边界、外部的环境变迁。这些“外部”被各理论当作自明的前提，而不是需要被解释的产物。正因为每一个理论都站在“系统/外部”这一二元区分的地基上，它们才得以建立起各自的分析框架：免疫是对“外部”他者的排斥；防御是对“外部”威胁的回避；置换是偏离了“外部”给定的原初目标；自毁是把“自身”错当成了“外部”敌人；陷阱是由“外部”环境的变迁所触发。

然而，本文在分析中所逐渐揭示的，恰是这样一个过程：一个规范性系统在长期维持自身秩序的过程中，通过精细化的合规、系统化的审核、制度化的培训，逐步将其内部所有真实的外部参照物——那些无法被既有规范消解的真实问题、真实冲突、真实的认知困难——一一消解、吸纳或绕过。当这个过程进行到担保虚化的晚期阶段时，系统面临一个奇怪的困境：它已经太成功地将“外部”消化掉了，以至于它丧失了辨识“维持”与“僵死”之间边界的参照系。此时，系统为了继续向自己证明自己仍在“运转”而非“空转”，不是去重新打开一个真正的外部，而是开始自己生产出一种看似在“外部”的批判位置——设立“独立挑刺人”、雇佣“颠覆性创新顾问”、举办“打破既有思维的头脑风暴会”——以填补外部缺席留下的虚空。

本文将此现象命名为“假性外部生产”，并主张这一机制是理解当前AI时代知识主体性危机的关键：并非“AI取代了人的判断”，而是相关认知系统在通过一套精密的规范性维持（高质量的AI辅助、标准化的批判性思维培训、精细的学术写作辅导）将判断的困难节点逐一平滑掉之后，开始将“对AI局限性的批判”本身生产为一个可被管理的、可被纳入既有规范轨道的“外部声音”，以此来证明“我们并没有被AI掏空，我们仍在批判它”。而这个被生产出来的批判位置，恰是规范性自噬的最高形态：它不是病理的出口，而是病理为了续命而为自己制造的必须的幻肢。

这一判断是本文对既有文献的最核心贡献。它指出了五个竞争概念共同站在其上的那块地基——系统拥有一个可被识别的外部——本身在特定条件下就是系统为掩盖内部虚空而生产出来的产物。这一反转将分析的焦点从“系统如何应对外部”转向“系统如何生产出它声称要回应的那个外部本身”，从而在既有学术库存中打开了一个集体不可见的研究领域。第五节的分析与讨论将逐层展开这一论证。

3 理论框架

基于上述文献对话中锁定的理论创新空间，本节正式构建“规范性自噬”的理论框架。该框架由三个核心概念——互锁式认知空转、排斥形态的无声交替、显影的自噬——以及一个升维概念——假性外部生产——组成，它们共同构成一个具有清晰时序形态和可证伪条件的理论整体。

3.1 核心概念的操作性定义

在本文中，“规范性系统”被界定为任何一个以“做出符合特定标准的正确判断”为核心功能目标、并为此建立了维持性秩序（如审核流程、分级标准、培训制度、同行评议）的社会认知系统。这一界定不限于正式组织（如大学、研究机构、专业协会），也包括弥漫性的认知共同体（如学科领域、学术流派、专业人士网络）和广义的教育体系。

“判断主体性”被操作性地定义为：一个认知者在不依赖标准化模板或外部权威的替代性判断的情况下，面对模糊、复杂、信息不完备的真实问题，能够经由亲身经历的认知张力（困惑、试错、反复推敲、自我推翻），做出一个可被自己负责、可被他人质疑、并能够经得起跨情境检验的判断的能力。这一概念区分于“决策能力”或“问题解决技能”——后者通常可以通过训练获得并可在标准化测试中衡量，而前者强调的是一种对判断过程本身的拥有感和担保能力：认知者能够说出“这是我的判断，我基于以下理由为之负责”，并且能够在他人的有力反驳面前，不简单地退回到“这是标准流程要求的”或“这是AI建议的”这类外部担保。

“劳动基底”被界定为：认知者在做出负责任判断的过程中，所必须亲历的那些无法被外包、无法被压缩、且其消耗本身就是判断生成之必要条件的认知活动——包括在困惑中反复推敲、在与他人的摔跤式对质中被迫澄清自己的预设、在信息不完备时承受做出抉择的压力、以及在自己的判断被证伪后亲历那种“我的地基裂了”的震荡。这一概念之所以称为“劳动”基底，是强调这不是一种静态的“知识储备”或“技能清单”，而是一种在动态使用中不断被消耗、又通过在真实判断中的反复使用而不断被再生的活性的认知能力。

“担保虚化”被界定为：一种系统性的过程，通过这一过程，认知者做出负责任判断所需的劳动基底，被日益精细化的合规流程、审查节点、模板化培训和AI辅助系统所逐步替代绕过，以至于个体虽然仍能在形式上“做出判断”（即在若干选项中做出选择并附上格式正确的理由陈述），但其所做出的判断不再携带着那种唯有亲历认知断层才能获得的、内生于判断者自身的认知重量。判断的担保从个体内在的认知劳动过程，转移到了外部合规流程的正确执行之上。本文将此称为“担保的虚化”——不是担保消失了，而是担保的实质内核（个体的负责任的判断劳动）被替换为担保的空壳（流程合规的形式正确性）。

“显影的自噬”被界定为：一种二阶的、反思性的病理机制，通过这一机制，以揭示和诊断上述担保虚化及其影响为目标的批判性活动（包括学术批判、诊断性教育、自我反思实践），因其在传播和教学过程中将认知断层的原初体验压缩为可被传授的抽象知识（“这里有坑，注意避开”），而在更大范围内稀释了后续认知者亲历那些断层的体验，从而使诊断行为本身无意中成为它所诊断的病理的新一轮执行者。关键在于，“显影”——将隐性规则显性化、将默会知识编码化——这一在这个过程中被普遍视为进步和赋权的行为，恰是担保虚化得以在更大规模上被复制的最有效方式。

“假性外部生产”被界定为：当一个规范性系统在担保虚化和显影自噬的长期作用下，其内部所有真实的外部参照物（真正的异见、不可被既有规范消解的真实问题、来自系统边界之外的冲击）都已被充分消化或隔离之后，系统为继续维持“我仍在回应外部挑战”这一自我确证的功能，而主动生产出一种看似在系统“外部”的批判位置的过程。这一被生产出的批判位置，其形式边界和讨论议程都在系统的维持性规范框架内被预先划定，因此它不会对系统构成真正的威胁；它的功能不是促进系统的改变，而是为系统的继续维持提供“我仍在自我更新”的自我叙事所需的参照物。

3.2 三阶段机制的命题体系

基于上述核心概念的操作性定义，本文提出以下相互咬合的命题，构成规范性自噬的理论主干。

命题1：互锁式认知空转。在一个高度规范化的认知系统中，当成员试图提出一个无法被既有概念网络充分归类的判断时（即遭遇“语言地形空白”），系统的维持性免疫应答——审核质疑、默会压力、对“不符合学科规范”的标示——并不直接压制该判断的提出，而是通过触发“请用已有术语定义你的新想法”的合法性追问，消耗提出者的“概念创生力”。提出者被迫花费大量精力在证明“我的问题是一个真问题”上，而不是在推进问题本身的解决上。与此同时，这一消耗过程产生的“降温效应”，使后续其他潜在的突破者在入水之前就已收脚——不是害怕被惩罚，而是预判了提出新问题所需的前置辩护成本过高。由此，规范性的维持动作（要求新判断必须能被现有术语体系正当化）与认知惰性（因为术语空白而无法正当化的新判断无法被认真对待），在语言地形空白处形成互锁，系统性地阻断了认知边界的突破。

命题2：排斥形态的无声交替。上述互锁式认知空转在系统的早期可能表现为显性的“惩罚异见者”——提出不符合主流范式的判断会受到制度性的阻力，从拒稿、拒评到学术边缘化。然而，随着系统规范化程度的提高，排斥的形态逐渐从“惩罚异见”演变为“使异见无从产生”。这种演进的发生路径是：日益精细化的合规流程、模板化的判断格式、多节点的审核制度、以及“最佳实践”的通用培训，在成员做出判断的每一个步骤上，都预先提供了一套“正确的方式”——如何提问、如何论证、如何引用、如何回应反对意见。这套方式并非不让成员思考，而是使得按照这套方式进行的思考，天然地倾向于产出那些在既有规范框架内被认为是“合格判断”的产物。个体不再需要感受到被压制的威胁，因为他们经由这套方式做出的判断，从源头起就已经是“合规”的——那些可能冒犯规范的粗糙洞见，在还没有清晰到能够被拿到台面上讨论的程度之前，就在这套“正确方式”的研磨下消失了。排斥从一种“对已产生的异见的惩罚”，演变为一种“使异见丧失产生条件”的基底掏空。本文将此称为从“急性排斥”到“慢性衰竭”的排斥形态的交替。

命题3：显影的自噬。针对上述病理的诊断和批判行为——无论是用批判性思维训练来让学生警惕“算法的偏见”，还是用学术研讨会来揭露“既有范式的压制”——在将其诊断成果传播给更多认知者的过程中，不可避免地面临一个困境：为了使其诊断可被传授和习得，诊断者必须将认知断层的原初体验——那些使人困惑、挣扎、反复推翻自己的痛苦过程——压缩为可被语言传递的抽象的“陷阱地图”。这种压缩在传播效率上是高效的，但它在每一次传播中，都会稀释接受者那一端亲历原初断层的体验强度。断层在被标注“此处有坑”的同一刻，对后来的认知者而言，就从一个“必须自己摸索才能通过”的困难节点，变成了一个“已经被前人探明、只需按导航绕行”的路标。由此，诊断行为在提升系统对既有病理的免疫意识的同时，恰恰衰减了系统成员在未来遭遇新断层时进行原初探

索所需的认知肌肉——那种在无人导航的困境中自己摸索出路的能力。这一自噬闭环堵死了“通过更透彻的诊断来拯救系统”的简单出路。

命题4：假性外部生产。当上述三个阶段在时间中递次展开——互锁式空转消耗了概念创生力，排斥形态交替掏空了判断的担保基底，显影自噬又使连诊断这一抵抗行为都被病理逻辑部分收编——系统最终进入一个悖论性状态：它已经高度成功地将所有真实的内部批判和不可消化的外部冲击都加工成了合规轨道内的“已知问题”，以至于维持自身运转所需的“我们仍在回应挑战”的自我叙事，也因缺乏真实的挑战素材而难以继续。在这一状态下，系统开始主动诱导、生产出一种看似在“外部”的批判位置——例如，设立“创新挑战者”角色、定期邀请“外部尖锐批评者”来演讲、在组织内部制度化地鼓励“唱反调”——但这些“外部”批判的形式边界和议程设定，本身就是系统为维持运转而划定的。它们的真正功能不是带来不可预测的改变，而是为系统提供“我仍在被挑战、因而我仍然活着”的本体性确证。这不是一个阴谋论式的“上层设计”，而是一种弥漫性的系统需求：一个已经内在地耗尽了判断主体性的系统，需要借助外部批判的“幻肢痛”来维持自己还有肢体的错觉。

3.3 跨学科理论资源的有效调用边界

本框架在构建过程中调用了免疫学、认知科学和组织理论的跨学科资源。在此必须显式说明其调用有效的条件与边界，以避免装饰性引用。

免疫学概念（自身免疫、免疫应答、免疫记忆）在本文中的使用，是作为结构性的形式隐喻，而非实质性的因果解释。调用有效的边界是：规范性系统的维持性排斥动作，在功能动态上（而非在生物学机制上）展现出与生物免疫系统在某些抽象形式属性上的同构——包括对“非我”的识别、记忆效应、以及在特定条件下攻击自身功能性组织。这一隐喻的启发性在于，它为本框架提供了区分“急性排斥”与“慢性衰竭”的形态学概念工具。它失效的边界是：当分析需要涉入具体的生物学因果机制（如抗体-抗原结合、细胞因子级联）时，该隐喻立即失效，不能用于指导可操作的干预方案。本文在整个论证中，从未将任何经验主张建立在生物学因果链之上；所有可操作变量和证伪条件的设计，均锚定在社会认知的制度性变量之上。

认知卸载研究为本文提供了更谨慎的经验背景。外部工具能够降低工作记忆负荷并改善即时任务表现，人们也会根据外部信息是否可再次获得而调整记忆策略（Sparrow et al., 2011; Risko & Gilbert, 2016）。但这些研究并不能直接证明一切工具使用都会导致长期能力衰退；卸载的后果取决于任务、反馈、练习机会和保留的核验责任。本文据此提出一个新的、尚待检验的问题：当被压缩的不只是信息存储，而是问题界定、证据冲突处理和反复修订等判断工序时，长期结果是否会表现为独立判断质量下降。这个推论必须通过纵向研究或准实验验证，不能由现有认知卸载研究直接推出。

组织理论中的制度逻辑和组织学习传统，为本框架提供了经验层面的材料锚定和分析工具。尤其是关于“制度复杂性”和“组织衰变”的研究，为本文的“担保虚化”在社会组织层面的表现提供了大量可观察的案例形式。此调用的有效边界是：本文不讨论组织的正式结构变量（如层级设计、部门分工），也不宣称所有类型的组织都必然经历本文描述的自噬过程。本文的理论范围明确限定在那些以“判断的正确性”为核心合法化基础的系统内，而非一般的科层组织或企业。

4 方法论

4.1 研究设计

本文采用概念分析（conceptual analysis）的方法论路径，这是一种在哲学、社会理论和组织研究中被广泛使用的方法，其核心任务是对一个尚未被充分理论化的经验现象，通过系统化的概念建构、与既有理论的对话、以及可证伪条件的给出，提供一种新的理论命名和解释框架。本文并不依赖一手经验数据的收集，而是通过对既有文献和理论资源的系统整合与批判性重新配置，建构出一个具有清晰时序形态和可操作检验方向的理论框架。

这一方法论选择是基于研究对象的本体论特性。本文所关注的“判断主体性的系统性耗竭”现象，其核心机制（如担保虚化、显影的自噬）发生在认知劳动的基底层次，难以通过单一的截面量化研究或孤立的个案研究直接被“捕捉”——因为这些机制涉及的是长期性的、弥漫性的、且往往与合法性的维护行为交织在一起的过程，参与者自身可能在担保虚化的晚期阶段已经丧失了辨识自身状态的能力（正如免疫记忆被稀释后，病人不再记得当初的病症之痛）。因此，本文选择在概念和分析的层面先完成理论框架的建构——为后续的实证研究提供可操作的核心概念和可证伪的假设方向——然后再由未来的实证研究在具体领域中检验、修正或推翻该框架。

本文采用机制型中层理论的建构路径：先限定适用对象，再提出变量关系、竞争解释与可能推翻理论的观察结果。它既不同于追求普遍有效的宏大理论，也不同于只描述个案的解释性叙事。组织学习研究已经揭示了规则、惯例、历史依赖以及探索—利用失衡如何塑造组织适应（Levitt & March, 1988; March, 1991）；本文在此基础上进一步追问，维护判断质量的程序在什么条件下会削弱判断形成过程。因此，“规范性自噬”应被视为一个等待比较研究检验的机制假设，而不是凭借概念自治即可成立的封闭体系。

4.2 概念建构的逻辑步骤

本文的概念建构遵循以下步骤：

第一步，问题现象的识别与命名。本文从当前AI时代知识生产领域广泛讨论但未被精确诊断的经验现象——“依赖AI后判断力‘失重’”——出发，指出这一现象背后存在着一个比“技能退化”更深层的机制痕迹。

第二步，与最近邻概念的系统对话。本文在文献综述中逐一与五个最相关的既有概念（德里达的自体免疫、阿吉里斯的防御性例行、默顿的目标置换、医学自身免疫、能力陷阱）进行交锋，逐一核验是否每一个竞争概念在与本文所揭示的现象进行一对一对质时，都存在一个可被清晰指认的“它看不见什么”的盲区。这一“盲区交集”方法论（而非“文献缺口”方法论）确保了本文的理论创新不是既有概念的语境替换，也不是“还不曾有人研究”的空地宣称。

第三步，核心机制的建构与删件测试。本文提出的四个机制（互锁式认知空转、排斥形态交替、显影自噬、假性外部生产），在理论建构阶段经过了一个严格的“删件测试”：如果删除其中任何一个机制，整个理论是否仍能在另外三个机制的支撑下独立站立？只有在满足“删除任何一个机制，整

个理论就会塌陷在某一个关键环节”的条件下，该机制才被认定为整个理论不可卸载的构成部分。本文在最终的典范建构中给出的四个机制，每一个的删除后果都已在文本中被显式论证。

第四步，证伪条件的设计。本文为上述理论框架给出了一个能够排除最强竞争解释的证伪检验方案（见5.4节）。该方案的操作性设计确保了本文不是在宣称一个无法被事实推翻的循环论证，而是在提供一个可能被未来的实证研究证伪的经验假设。

4.3 方法局限

本文的概念分析方法论存在以下局限，需预先诚实陈述。

第一，本文的理论建构主要基于逻辑推演和跨学科的形式映射，而非一手经验材料的实证支撑。虽然本文在分析过程中附着了多个可观察的经验场景作为例示（如教育、企业治理、AI辅助认知），但这些场景并未经过系统的案例研究或数据收集。本文当前的贡献在于提供了一个可被经验检验的理论框架和操作化的概念工具，而非已经完成了对这一框架的经验验证。

第二，本文在调用跨学科类比（尤其是免疫学隐喻）时，虽然已显式标注了类比的有效边界，但仍面临一个固有风险：隐喻可能在不经意间引导论证方向，使得某些本该被严肃考虑的经验细节因为“不符合隐喻的逻辑”而被忽略。为减小此风险，本文在证伪设计中将理论的可检验性锚定在非隐喻性的操作化变量（合规流程的颗粒度、审核节点数、异质信号识别力测试）之上，以确保理论最终的生死取决于经验证据，而非隐喻的修辞力量。

第三，本文所界定的核心概念（尤其是“判断主体性”和“担保虚化”）的操作化难度较高，可能在未来的实证研究中面临测量效度的挑战。本文对此的策略是：在概念建构阶段保持界定的精确性，而将具体的测量工具开发（如判断主体性的量尺设计、担保虚化的代理指标选择）留给后续的方法论专项研究。本文在证伪方案中已提供了可供进一步操作化的代理变量方向。

5 分析与讨论

5.1 规范性自噬的触发：互锁式认知空转

让我们将一个典型的知识工作者放在分析的起点。她是一名在某个高度专业化的学科领域内接受过系统训练的研究者，掌握了该领域的主流理论框架、核心方法和被一致认可的经典问题。当她面对一个无法被该领域的既有概念网络充分归类的现象时，她遭遇的不仅仅是认知上的困惑，更是一种规范性的困境：她必须用学科已有的术语为她的新问题提供合法性论证，但这些术语本身正是在一个不曾涵盖该现象的历史语境中形成的。她用在一个旧语境中准确但在新现象面前错位的词来描述问题，被同行指出“概念使用不当”；她试图铸造一个新术语，被质疑“不符合学科规范”；她试图绕开既有框架直接描述现象，被批评“缺乏理论根基”。

这就是本文所称的“语言地形空白”：既有概念网络的覆盖范围与真实问题地形之间出现的裂缝。在这道裂缝中，研究者遭遇的不是对错误判断的否定（她的判断可能在大方向上是对的），而是她的判断被悬置在“无法被认真对待”的中间状态——因为认真对待一个判断的前提，是该判断能够被共同体的既有语言工具所“接住”。如果研究者坚持在这道裂缝中工作，她就必须额外支付一

笔“概念创生力”的成本：她不仅要解决现象本身的问题，还要花费大量精力在铸造能够承载这个问题的语言容器上。而这笔成本在系统的规范性审核中是不被看见的——因为在她成功铸造出足够好的新术语之前，系统看不到“有什么问题需要被铸造术语”。

这笔被持续消耗的“概念创生力”是规范系统维持自身秩序时产生的第一种代谢废物。它的消耗不会直接导致系统的崩溃——恰恰相反，系统在短期内会因为它有效过滤掉了大量“还不够成熟”的新想法而显得更为高效、更有秩序。但长期来看，这种消耗的积累效应是：共同体内逐渐减少的，不是已经成型的异见数量，而是能够产生异见的那种“在语言空白区滞留、摸索、反复试错”的认知耐力。当下一批研究者进入该领域时，他们接受的训练就是建立在那张已被前人标注过的概念地图之上的——他们在还没有亲历地形空白之前，就被告知了地图上哪些区域是“尚未被研究”的。这种标注本身，在节省了他们探索时间的同时，也使他们丧失了空白中铸造新词的认知肌肉记忆。

这一机制解释了为什么那些“突然”被颠覆的学科典范，往往不是在它们面临最多外部批评的时候崩溃的，而是在它们最成功地将所有可被批评的点都标注为“已知的局限”之后的某个不可预见的时刻脆断的——因为标注已知局限的过程，已经把共同体成员识别未知局限所需的认知感受力磨平了。

5.2 从急性排斥到慢性衰竭：排斥形态的无声交替

互锁式认知空转描述了系统在语言空白处对突破性判断的“初筛”机制。在系统的早期，这一初筛常常以一种“有迹可循”的排斥方式操作：提出不合主流的新判断的研究者，会经历可被第三人观测的压制痕迹——被拒稿、被拒评、被边缘化于学术会议之外、被有权势的导师劝告“先做成熟的方向”。这些痕迹是显性的，它们使得外人可以绘制出一张“学科压制的地图”，也可以使抵抗者识别出“谁是压制者”。这是排斥的急性阶段。

然而，随着系统的规范化程度的提高——这通常被参与者体验为“学术质量的提升”和“学科建设的完善”——这种急性排斥逐渐被一种更精致、也更不可见的排斥形态所取代。日益完善的审稿标准不再需要审稿人明确地说“你的观点不主流，我拒你”；审稿人可以诚实地说“你的论证没有达到本刊要求的方法论标准”，而这个“方法论标准”本身就是由该领域的主流范式的维护者们在长年累月中定义和细化的。它不是专门为了压制异见而设计的，但当它被运用到一个站在范式之外的新判断上时，它自然就会产生压制效应——因为该判断所依靠的方法论基础，很可能恰恰就是它在挑战的那个范式所定义的“方法”之外的某种尚未被规范化的探索方式。

在更深的层次上，精细化的培训制度和“最佳实践”的知识传递，在帮助新人快速上手的同时，也在系统性地压缩他们自己摸索出“不同于最佳实践的工作方式”的可能性。一个博士生在入学之后，被系统地训练用本学科的规范方式——写文献综述、构建理论框架、选择研究方法、呈现数据分析——来处理他的研究对象。这套方式在大部分情况下是有用的，它帮助他避免了许多初学者常见的错误。但它在成功的同时，也在做一件不被注意的事：它把他可能提出的那些粗糙的、不成形的、暂时还无法用这套规范方式“合法化”的直觉判断，在他自己意识到它们可能具有价值之前，就已经温柔地拦在了门槛之外。他不是被禁止了提问，而是他还没有学会如何把他的问题用一种可以被认真对待的方式问出来，而当他学会了那种提问方式之后，他的问题可能已经不再是当初那个捅破范式的问题了。

这就是本文所称的“担保虚化”：系统通过日益精细的维持性动作，不再需要惩罚异见者，因为它在更早的阶段——在潜在异见者还在学习如何思考的阶段——就已经将他们的思考引导到了合规的轨道内。被排斥的不再是“已产出的异见”，而是“产出异见的劳动条件”。这一排斥的形态转变，使得系统的免疫性自噬从一种可以被观测和反抗的“压迫”，演变成一种弥漫性的、无法被归因于具体压迫者的“衰竭”。在衰竭状态下，个体不再感到自己“被压制了”，他们只是诚实地感觉到“没有足够好的新想法”——而这个诚实的感觉本身，就是担保已经被虚化、判断基底已经被掏空后最准确的症状。

5.3 诊断的悖论：显影的自噬与假性外部生产

现在我们来到整个论证最反直觉的一环。如果上述两个机制的诊断是正确的——如果一个高度规范化的认知系统正在通过排斥形态的交替，系统性地消耗其成员做出负责任判断的劳动基底——那么我们显然应该去做一件事：揭示它。我们应该用研究来曝光隐性排斥，我们应该在教育中训练学生识别“范式压制”的痕迹，我们应该用AI这面镜子来照出规范系统内部的断层，让所有人都看见它们。这正是当前大量关于“批判性思维”、“算法透明”、“认知偏见”的论述所呼吁的。本文的作者也曾在这条路上怀抱着真诚的热情。

然而，本文在此要提出的，是一个令人不安的二级效应。当我们将那些隐性的认知断层——那些让前人在困惑中挣扎、在无路标的地形中摸索出道路的困难节点——用精准的语言显影、标注、编入教材、嵌入AI辅助系统时，我们在做一件双重的事情。一方面，我们确实使后来者避免了前人在同一个坑里的无谓挣扎，提升了认知效率。但另一方面，我们在每一次成功的显影中，都压缩了后来者亲历那个断层的原初体验。断层在被命名为“此处有坑”的那一瞬，对后来者而言，它就不再是“一个我必须自己摸索才能通过的困难节点”，而是“一个已经被前人探明、只需按导航绕行的路标”。

这里的机制不是“知道了陷阱就不需要再学”——这当然是正确的，也是所有教育的基础逻辑。这里的机制在于：那种在无人导航的困境中自己摸索出路的认知能力，并不是一种可以在绕行陷阱的同时独立发育的、可被隔离训练的技能。它是在亲历一个又一个不可预见、未被标注的原始断层中，通过反复承受“我的既有判断在这里失效了”的冲击，而逐渐生长出来的一种对于复杂性的感受力和在模糊中做出负责任抉择的胆量。当教育系统和AI辅助系统越来越成功地提前标注了越来越多断层之后，后来的认知者在他们受训的整个过程中，所经历的“既有判断突然失效”的体验的总量和强度，就发生了系统性的衰减。于是他们毕业时，拥有了比前代更精确的错误地图，却可能失去了在没有地图的地形中摸路的原始能力。

这构成了“诊断知识的经验压缩”风险：病理一旦被清晰命名，后来者就更容易获得路线图，却可能更少经历发现问题、形成概念和修正错误的完整过程。这里不能推出“诊断越成功，能力必然越弱”。更可检验的预测是：在教学内容相同的条件下，只接受陷阱清单与标准答案的学习者，可能不如先独立处理陌生问题、再接受诊断反馈的学习者，能够把判断策略迁移到未被标注的新情境。若两组在迁移任务上的表现没有稳定差异，本文关于经验压缩的核心机制就应被削弱或放弃。

这一机制将本文的论证推入了一个更深的困境。如果显影的自噬是对的，那么“通过更透彻的诊断来拯救系统”这条路上就竖起了一道无法绕过的墙：诊断本身就是新一轮担保虚化的高效执行者。但这是否意味着本文的整个论证只是一个精致的虚无主义——我们在诊断一个不可治愈的疾病，且在诊断的同时强化着疾病？本文的回答是否定的，但这需要一个更艰难的概念跃迁：我们必须区分“真

实的诊断”与“假性诊断”，而这一区分的可能性，取决于诊断者是否意识到自己对“外部于系统的批判位置”的占有本身就是系统为维持运转而生产的假性外部。

当一个系统已经进入了担保虚化完成的晚期阶段，所有真实的、不可被既有规范消化的内部批判和外部冲击都已沉寂之后，系统面临一个存在性危机：它失去了辨识“我仍在运转”与“我已在空转”之区别的参照物。此时，它最深的恐惧不是“被批判攻破”，而是“连批判它的人都不再出现了”——因为那意味着系统已经彻底丧失了外部，沦为一种没有任何外在参照的、绝对自我指涉的空转。为防止这一寂灭，系统会开始主动生产出一种看似在“外部”的批判位置：设立“独立挑刺人”角色、邀请“外部尖锐批评者”来发表演讲、在组织内部制度化地鼓励“唱反调”的成员。但这一“外部”批判的议程设定、讨论边界和可被允许的激进程度，本身就在系统的既有规范框架内被预先划定；它的真正功能不是带来不可预测的改变，而是为系统提供一种“我仍在被挑战，因而我仍然活着”的本体性慰藉。

当认知系统的批判者们——包括本文的作者——在系统地揭示“AI压缩了判断劳动”、“既有范式压制了创新”、“合规流程在掏空主体性”等等的时候，如果他们的批判工作满足两个条件，则其本身就落在假性外部生产的逻辑之内：（1）批判者的批判语言和概念框架，来自于并一直停留于该认知系统所提供和认可的那些“批判性话语传统”（如“技术批判”、“算法偏见”、“范式革命”等等）的内部，从未触及自身的劳动基底被担保虚化这一更根本的病理；（2）批判者及其受众之间的互动，本身构成了一种“我们正在负责任地反思自身”的仪式性循环——这种循环使系统内的成员在体验了一次“深刻的自我批判”之后，感觉到“我们已经意识到问题了，这本身就是进步的证明”，从而反而放松了在自身日常判断中去亲历认知断层的必要性。

真正的诊断——如果它可能存在——不在于为系统提供一套更精致的批判话语，而在于拒绝占有系统为维持运转而预留的那个“批判者位置”。这意味着，真实的诊断者必须不断地在自身的工作中接受“我的这个判断，是否也在担保虚化的逻辑之中？”的拷问，并时刻准备承认“是的，它可能也是”这一可能的答案。同时，他们必须将诊断的终点不是安放在“我们揭破了病理”的满足感上，而是安放在一个更朴素、也更不可被收编的地方：自己去面对一个真实问题，耗尽自己的劳动基底，做出一个负责的、不可被合规流程替代的、愿意为之承担被推翻风险的判断。这个判断可能是不完美的，可能是有盲区的，可能在未来被证伪——但它有一个虚假批判永远无法拥有的特征：它携带着判断者亲历认知断层的全部重量，因而它是“活”的。

5.4 可证伪检验设计：排除环境变迁的竞争解释

为了确保本文的理论框架不是一个无法被事实推翻的封闭叙述，本节为规范性自噬的核心机制提供一个可操作的证伪检验方案。该方案聚焦于理论中最要害的那个预测：维持性动作的精细化本身，在环境高度稳定的条件下，仍能内生性地消耗判断主体性。这一预测将规范性自噬与能力陷阱（competency trap）这一最强竞争解释在最关键的区分点上推入可检验的对质。

最强竞争解释。本文所面临的最朴素、也最有力的替代解释是：那些规范性系统之所以出现判断力的衰退，不是因为其维持性动作自身有什么内在的自噬效应，而纯粹是因为它们所处的环境发生了变迁——旧有的判断模式不能适应新的外部条件，而系统由于惯性未能及时调整。这一解释将衰竭的根源完全定位在“系统-环境”的不匹配上，而非在维持性机制自身的消耗效应上。如果这一解释成

立，那么在环境高度稳定、不存在“适应失败”的条件下，维持性动作的精细化就不应该系统地导致判断力衰退——因为并不存在一个需要适应但未被适应的“新环境”来惩罚旧判断的失效。

检验设计。采用同质多案例比较设计（most-similar-case comparative design）。选取至少六家处于同一高度稳定行业中的在位组织（如特定传统制造业、基础公共事业、不可替代的专业服务领域），控制其在过去十年内面临的外部环境变迁幅度——技术轨迹无明显跃迁、市场需求无结构性转变、监管政策框架无重大修改。进一步控制其总体规模、历史声誉、研发投入占营收比例和过往的创新记录。在此条件下，测量两个核心变量的组织间差异：

- 自变量“维持性动作的精确度”：可操作化为内部合规流程的颗粒度（需经审核的决策节点数与层级）、模板化操作程序在所有常规任务中的覆盖率、成功提案所需的审核步骤数、以及组织成员在做出专业判断时被要求填写或参照的标准化文件的平均数量。

- 因变量“内部判断力”：操作化为组织对一项预设的、来自于本行业之外但在未来五至十年内具有潜在颠覆性影响的技术趋势的早期识别与正确评估能力。该评估可通过独立的专家裁判组（盲于各组织的自变量得分）对组织内部的相关技术报告、战略文件、或关键成员在结构化访谈中所展示的“对异质信号的敏感度和正确辨识力”进行独立评分。

证伪条件。若在控制了环境和组织背景变量的前提下，维持性动作的精确度与内部判断力之间不存在有统计显著性的负向剂量-反应关系——即流程最精细的组织并未比流程最精简的组织表现出更低的异质信号识别力——则本理论的核心预测为伪。此时，我们可以推断规范性系统的判断力衰退主要是由环境变迁触发的能力陷阱效应，而非由维持性动作本身的内生消耗。若该负向关系存在且呈剂量-反应模式——即维持性动作精确度每增加一个单位，内部判断力呈逐步递减趋势——则能力陷阱解释在此环境中无法成立（因为环境被控制为稳定，排除环境变迁干扰），而本理论的独特机制获得了经验支持。

若被证伪，将学到。如果上述证伪条件成立（即无显著负向关系），规范性自噬的理论就需要被退回到一个更受限的范围：它可能只在环境发生变迁时才显著运转（此时它与能力陷阱共享解释域），或者在特定类型的规范性系统中才成立（如学术学科而非企业组织）。这将帮助划定本理论的有效边界，避免其被过度泛化为一个“解释一切规范系统衰败”的万能框架。如果证伪条件不成立（即存在负向剂量-反应关系），规范性自噬则获得了一个独立于能力陷阱的经验证据，其干预意涵也随之改变：对于高度稳定的行业或领域而言，防范衰败的关键不在于“紧跟环境变化”，而在于从内部设计上为维持性动作设定一个自我设限——在维持必要规范性的同时，故意保留一些“合规流程不允许绕过”的粗糙判断区域，让成员在其中不得不自己摔跤。

5.5 综合讨论：理论贡献与干预意涵

整合以上分析，本文的理论贡献可以凝缩为以下三点。

第一，本文在五个高度相关的既有概念（德里达的自体免疫、阿吉里斯的防御性例行、默顿的目标置换、医学自身免疫、能力陷阱）所共同占据的“系统/外部”二元地基之外，命名了一个因为它们共同站在那块地基上而集体不可见的原始现象：系统在担保虚化的晚期阶段，会主动生产出外部批判位置以填补内部虚空。这一“假性外部生产”的识别，将分析的焦点从“系统如何回应外部批判”转

向“系统如何生产出它声称要回应的那个外部本身”，从而为规范系统的衰变研究打开了一个新的问题域。

第二，本文为这一现象赋予了具有清晰时序形态和可操作判据的理论框架——规范性自噬的三阶段机制。从互锁式认知空转（触发和初期消耗），经由排斥形态的无声交替（从急性压迫到慢性衰竭的演变），到显影的自噬（诊断被病理收编的反身闭环），再到假性外部生产（虚空自我填充的最终形态），这一过程不是四个并列的“方面”，而是一个有先后次序、每一阶段都依赖前一阶段为其创造条件的生成机制序列。这一时序形态赋予框架一种动态的、而非静态的解释力。

第三，本文通过给出可排除最强竞争解释的可证伪检验方案，将这一理论框架拖出哲学思辨的安全区，推入可被经验事实否证的领域。这一可证伪性——而非仅仅的“凝练”或“深刻”——是本文所追求的学术品格的核心所在。

在实践干预层面，本文的理论引出了一个令人不安但必须被认真对待的启发：那些被普遍视为“对抗AI侵蚀判断力”的策略——强化批判性思维教育、揭示算法偏见、培训“人工智能素养”——如果在设计上不审慎地考虑显影的自噬效应，可能在减轻某些症状的同时，加深了疾病的根本病理。当我们将“如何不被AI掏空判断力”本身变成一个可被课程化、可被AI辅助传授的“技能模块”时，我们是在用担保虚化的方式来对抗担保虚化——其结果不是治愈疾病，而是让患者在服用更精致的安慰剂后，更加确信自己正在康复。

本文因此倾向于一个更朴素、更不令人满意的干预方向：不是去设计一套“反自噬的最佳实践流程”——因为这种流程本身就将作为新一轮担保虚化的执行者——而是去在每个具体的认知劳动场景中，有意识地保留、甚至刻意制造一些“AI和合规流程都不允许帮你绕过去的”认知困难。这意味着，在课程设计中，不是让学生先用AI生成草稿再批判反思（这仍是AI先帮他们绕过了第一道最困难的原创生成），而是让他们先在没有AI辅助的条件下自己摔跤，在摔跤后在教师的帮助下一起反思“我刚才经历了什么”，然后再引入AI作为反思的对照——而非替代。在组织决策中，不是在流程末尾设置一个“挑战者角色”来形式化地质疑既有判断，而是在流程的最初阶段，要求提案者不使用任何标准化模板，用自己的语言写出他所面对的问题为什么是一个真问题、他的判断在什么条件下可能是错的。这些干预的共同核心是：不去试图根除认知困难，而是去保护认知困难不被过早地平滑掉。这种保护本身，就是对抗担保虚化的最原始、也最有效的方式。

6 结论

6.1 核心发现

本文论证了一个规范性系统——一个以做出正确判断为核心目标的社会认知系统——在其维持自身秩序的过程中，会启动一套最终反过来侵蚀其赖以做出判断的核心能力的免疫性机制。这一“规范性自噬”过程经历三个阶段：首先，当判断遭遇语言地形的空白区域时，规范性的维持应答会消耗认知者的概念创生力，在维持秩序的同时阻断认知边界的突破（互锁式认知空转）；随后，系统的排斥形态从可观测的“惩罚异见者”演变为不可见的“担保虚化”——日益精细化的合规流程系统性地掏空了个体做出负责任判断所需的劳动基底，使异见在产生条件上就被消解，而非在产生后被压制（排斥形态的无声交替）；最终，以揭示和诊断上述病理为目标的批判性行为，因其在传播中压缩了认知

断层的原初体验，而成为担保虚化在更大规模上的执行者（显影的自噬）。这三个机制的递次展开，将系统推入一个最终阶段：在内部判断主体性已然枯竭后，系统为继续维持“我仍在运转”的自我确证，开始主动生产出看似在“外部”的批判位置——这恰是规范性自噬的最高形态（假性外部生产）。

这一理论框架的贡献不在于发现了一个全新的孤立现象，而在于将若干分布在哲学、组织理论、社会学和免疫学中各自分离但高度相关的洞见，整合进一个具有时序形态、可操作变量和清晰可证伪条件的统一分析框架内。本文在德里达的自体免疫哲学结构、阿吉里斯的防御性例行、默顿的目标置换等经典概念所共同贡献的洞见的基础上，进一步命名了一个因为这些概念共享同一理论预设而集体不可见的原始现象：外部于系统的批判位置本身，在特定条件下就是系统以病理方式生产的产物，而非病理之外的拯救者。

6.2 理论贡献

本文的理论贡献有三重。

第一，概念贡献。本文提出的“担保虚化”、“显影的自噬”和“假性外部生产”三个新概念，为分析规范性系统自我侵蚀的不同维度提供了可被未来研究进一步操作化的理论工具。这些概念在每一个竞争概念面前都给出了可判定的差异条件，因此不是既有术语的语境替换，也不是对一个已经被命名过现象的重命名。

第二，机制贡献。本文为规范性系统的衰变提供了一个时序性的动态机制——从初期的互锁式空转，经中期的排斥形态交替，到晚期的显影自噬与假性外部生产——而非一个静态的“原因清单”。这一时序性赋予理论以预测力：它能够在系统当前所处的阶段上对后续的演变方向做出有经验风险方面的预测，例如预测一个担保虚化程度已高的系统，在引入外部批判机制时，将更可能生产出一个假性外部而非获得真正的自我更新。

第三，方法论贡献。本文通过“盲区交集”策略（而非传统的“文献缺口”策略）锁定理论创新的不可还原性——即论证所有最近竞争概念共同站在其上的那个未经追问的地基，然后将本文的理论定位在那块地基之下那个被集体忽略的现象域。这一策略为概念分析型研究提供了一种比“尚未有人研究”更坚固的、可与既有知识体系进行严格对话的创新路径。

6.3 实践与政策含义

作为一项概念分析型研究，本文不宣称有能力直接给出“如何避免规范性自噬”的操作手册。然而，本文的理论分析至少指向以下三个值得实践者认真对待的方向。

第一，对“最佳实践”和“标准化流程”的审慎反思。本文的分析提示，那些在设计之初以提高判断质量和降低错误率为目标的标准化流程，在越过某个临界点之后，可能开始内生产生一个未被预料到的副作用——掏空判断者经由亲历认知困难而获得的那种难以被模板化的判断品质。这要求组织在设计流程时，不是简单地追求“更精细的合规”或“更高质量的标准化”，而是在设计上刻意保留一些“流程管不到、必须自己判断”的粗糙区域。这些区域的保留，不应被理解为“流程还未完善”，而应被理解为“流程刻意为自己设定的边界”——正如免疫系统不是对所有外来物都发动攻击，它必须保留一个“不攻击”的空间，否则就会发生自身免疫病。

第二，对批判性思维教育和“AI素养”培训的重新审视。本文提出，将“对抗AI的侵蚀”本身变为一门可以被AI辅助教授和评估的课程，可能恰恰是担保虚化的新一轮执行者。本文的建议不是放弃这些教育努力，而是在教育设计的时序上给出一个具体的调整建议：先让学生在没有AI辅助的条件下自己面对真实认知困难并完成判断，然后再引入AI作为反思的对照；而不是先让AI生成一个答案，然后再让学生去“批判它”。前者的核心认知体验是“我在困难中摸索出判断”，后者的核心认知体验是“我在评判一个外部已经生成好的判断”——而这两种体验所调动的认知劳动基底的性质是完全不同的，前者生成了判断主体性，后者只是使用了一个更精致的判断客体的消费品。

第三，对学术评价体系的一个深层提醒。如果本文的分析成立，那么那些被评价为“高质量”和“规范化”的学术成果——它们通过了多轮严格的同行评议、严格遵循了方法论标准、有着无懈可击的论证结构——有可能恰恰是担保虚化的症状，而非担保健康的证明。因为担保虚化的最成功产物，正是那些“在形式上完美但不再携带判断者本人亲历认知断层的重量”的判断。本文并不天真地主张放弃同行评议和方法论标准，而是建议评价体系在单一的“规范性质量”维度之外，增加一个目前尚难以形式化但可被经验丰富的学者辨认出来的评价维度：这篇作品，是否读起来像是“作者自己在困难中摔出来”的判断，还是读起来像是“在正确的流程中组装出来”的判断。这一维度的引入，不是为了重新制造一个非此即彼的二元对立，而是为了在系统的免疫性规范过于强盛时，保留一个对它进行“免疫刹车”的调节入口。

6.4 研究局限

本文在以下方面存在局限，需诚实自陈。

第一，本文是一篇概念分析型研究，其理论机制（尤其是“显影的自噬”和“假性外部生产”）当前仍主要基于逻辑推演和跨学科隐喻的形式同构，尚未经过系统的实证检验。本文在5.4节给出的可证伪检验方案，尚停留在设计层面，未被实际执行。因此，本文的理论框架在当前阶段应被视为一个等待经验检验的、具有较高理论深度和跨情境解释力的假设体系，而非已完成验证的知识。

第二，本文在分析中将“规范性系统”作为一个抽象的总体来讨论，未深入区分不同类型的规范性系统（如学术学科、专业协会、企业内部研发组织、政府决策机构）之间可能存在的显著差异。这些差异可能调节自噬的发生速率、强度甚至方向，而本文的理论框架目前尚未提供对这些调节变量的系统性分析。未来的研究需要对不同类型规范性系统进行比较分析，以进一步精细化本文的理论边界。

第三，本文的“假性外部生产”机制，虽然在理论上有力地揭示了一个集体不可见的现象，但其经验识别面临固有的困难——如何区分“系统真正在回应外部批判”与“系统在生产一个假性外部来掩盖内部虚空”？本文目前尚未给出一个能被第三方观测者可靠应用的判据。5.4节的证伪设计锚定在更易操作化的“维持性动作-判断力衰退”剂量-反应关系上，未对假性外部生产本身提供单独的证伪检验。这是本文留给未来研究的一个核心难题。

6.5 未来研究方向

基于本文的理论框架和自我批评中识别出的局限，本文提出以下五个可生长的未来研究纲领。

研究纲领一：执行5.4节设计的同质多案例比较检验。招募一个跨学科研究团队，在至少六个处于同一高度稳定行业中的在位组织中进行变量测量和假设检验。若检验结果支持本文预测，则进一步探索维持性动作精确度与内部判断力衰退之间剂量-反应关系的具体函数形态（是线性、加速还是存在临界阈值）。若检验结果推翻本文预测，则转而探究规范性自噬的发生是否以某种环境变迁为必要条件，从而重新划定本文理论的解释边界。

研究纲领二：开发“担保虚化”的测量工具。基于本文的操作性定义，开发一套能够可靠区分“担保在个体判断劳动中内化”与“担保已虚化为外部合规流程的正确执行”的量尺或行为指标。该工具的开发需要认知科学与组织心理学的跨学科协作，可能结合判断质量评估、认知过程捕捉（如出声思维法）和纵向追踪设计（随时间检验同一批认知者在经历标准化流程培训后，其判断中的内部担保重量是否出现系统性的衰减）。

研究纲领三：设计并实施“干扰认知卸载”的准实验。在教育场景中，选取可比的课程班级，随机或准随机分配为两组：一组接收传统的高AI辅助+批判反思训练（先让AI生成，后让学生批判性分析）；另一组接收“障碍保留”式训练（先让学生在无AI下自己摔跤完成判断，然后在教师帮助下分析自身认知过程，最后引入AI作为对照反思工具）。在一个学期或更长时间后，比较两组学生在面对一个完全陌生的、超越课程既有内容覆盖范围的新问题时的独立判断质量。这一研究可首次为显影的自噬机制提供受控条件下的因果证据，而非仅仅的相关性观察。

研究纲领四：对“假性外部生产”进行深入的组织民族志研究。选取一到两个被其成员广泛认为“拥抱批判、欢迎外部挑战”的成熟组织（如以反思性著称的学术机构、或设有制度化“内部挑战者”角色的企业），进行长期参与观察和深度访谈。研究者需尤其关注：组织内被许可的“批判”的议程是如何设定的；批判的激进程度在何处触达不允许越过的边界；成员在参与这些批判实践后，是体验到了“系统正在发生实质变化”的确据感，还是体验到了“我们批判过了，然后一切照旧”的仪式性满足感，还是两种体验微妙地交替出现。这种民族志工作可以为假性外部生产的理论提供在单一量化研究中无法获得的、对机制微观运作的深描。

研究纲领五：跨学科边界的检验——“学科自噬比较研究”。将本文的理论框架应用于学术学科自身的演化研究。选取数个正处于不同历史阶段的学科领域（如一个成熟稳定但在近二十年内无重大理论突破的学科、一个当前正处于范式更替震荡期的学科、一个新兴但正在迅速建立起严格方法规训的学科），对它们的学术发表规范、方法论培训制度以及学领域内新理论创生的速率和来源（是产生于方法论最严格的中心机构，还是产生于方法论不那么严格的边缘机构？）进行系统的历史与比较分析。这一研究纲领既是对本文理论的跨情境检验，也可能为“学科自身的规范成熟与其理论创新能力之间是否存在倒U型关系”这一更为宏大的问题提供初步的回答。

附论·最近邻切割：规范性自噬不是德里达自体免疫、不是组织防御、不是目标置换

本文原稿已诚实列出三个对勘对象——德里达的自体免疫、阿吉里斯的组织防御、默顿的目标置换。这份自觉难得，但也把切割的责任摆到明处：一个主动承认自己站在三个成熟概念旁的框架，必

须证明它多切出了什么，否则就是三者的合写。此外还要处理一个新的近邻——它与本站金华《自饮性》同用「自噬/自我耗损」的隐喻，须一并划清。

第一个邻居是德里达的「自体免疫」（autoimmunity）——一个系统为自我保护而攻击自身保护机制的悖论。规范性自噬看起来就是它的组织版。切开的关键：自体免疫是一个哲学隐喻/结构描述，它指认悖论的存在，但不给出悖论发生的条件与程度；本文的「规范性自噬」是一个中层理论——它把这个悖论限定在一组可变的条件组合上（规范密度 × 任务新颖度 × 裁量空间 × 判断质量），主张只有在特定组合下自噬才发生，其余组合下规范是有益的。德里达说「系统可能自我攻击」，本文说「在规范密度高、任务新颖、裁量空间被压缩的特定象限里，维持性程序才会侵蚀判断经验，而在其他象限里不会」。可裁决分离：自体免疫作为纯结构描述不可证伪；规范性自噬预测一个具体的交互效应——判断质量的下降只在「高规范密度 × 高任务新颖度 × 低裁量空间」象限出现。若判断质量的下降与这三个变量的组合无关、或在所有象限均匀出现，则本文把德里达隐喻转成中层理论的努力失败。

第二个邻居是阿吉里斯的组织防御与「熟练的无能」。切开：组织防御的机制在个体的认知—情感层（Model I 行动理论使个体回避暴露性判断），它是一个心理—行为机制；规范性自噬的机制在制度的程序层（审核、方法训练、合规程序的加密本身压缩了探索经验），它不需要经过个体的防御心理即可发生——哪怕成员心态完全开放、毫无防御，只要程序密度高到压缩了探索空间，自噬照样发生。分离线可裁决：组织防御可由个体层面的干预（双环学习、心理安全感建设）缓解；规范性自噬不能靠改善个体心态缓解，因为侵蚀源于程序结构而非个体防御——这条线可在「成员心理安全感高但程序密度也高」的组织里检验：本文预测这类组织仍会出现判断经验的萎缩。

第三个邻居是默顿的「目标置换」（手段变成目的，官僚固守规则而忘记规则所服务的目标）。切开：目标置换讲的是价值理性向工具理性的错置（规则从手段异化为目的），它的落点在价值取向的漂移；规范性自噬讲的是判断官能的经验基础被程序抽干（不是「忘了目标」，而是「即使记得目标，也不再产生新判断所需的探索经验」）。一个组织完全可能时刻牢记目标、毫无目标置换，却仍因程序密度过高而陷入规范性自噬。可裁决：目标置换可由「重申目标、反官僚化」矫正；规范性自噬不能靠重申目标矫正，因为缺的不是目标意识，是探索经验的再生产条件。

第四，与本站金华《自饮性：制度化认知的成功何以系统性地耗损其认知能力》的划界。两文同用「自我耗损/自噬」的隐喻，但切的对象完全不同：金华的自饮性讲制度化预测系统在正常成功运转中消耗其赖以运转的公共认知基底（对象=预测/度量系统，机制=成功即消耗，近邻=古德哈特定律/操演性/公地悲剧）；本文的规范性自噬讲 AI 时代知识主体在高规范密度下丧失形成独立判断的探索经验（对象=认知劳动中的判断主体，机制=维持性程序压缩探索，近邻=自体免疫/组织防御/目标置换）。一个耗损的是系统的外部预测力，一个耗损的是主体的内部判断力；一个的解药在保留异质性，一个的解药在保留裁量空间。二文可互为参照：它们从预测系统与知识主体两端，共同指认了「制度化在成功中自我掏空」这一更大结构的两个侧面。

附论·可裁决预测：把哲学诊断压成有经验风险的交互效应

原稿已提出同质多案例比较、过程追踪、教育准实验三条检验路径。此处把核心命题收束为三条可被判错的预测。

预测一（象限特异性）：判断质量的下降是规范密度、任务新颖度、裁量空间三者的特定交互结果，集中出现在「高规范密度 × 高任务新颖度 × 低裁量空间」象限。若判断质量下降与该交互无关、或在所有象限均匀出现，则规范性自噬塌回一个不可证伪的普遍悲观。

预测二（绕过个体防御）：在成员心理安全感高、无明显组织防御的组织中，只要程序密度高且裁量空间被压缩，判断经验的萎缩仍然出现。若高心理安全感即可消除萎缩，则本文的「机制在程序层而非个体层」被证伪，规范性自噬还原为组织防御。

预测三（假性外部生产的可识别性）：「假性外部生产」预测——被制度化引入的批判，其议程、强度与后果被预先限定，因而不产生结构更新。可操作化：比较「自发涌现的批判」与「制度化安排的批判」在事后是否导致可测量的结构变动（规则修订、资源重配、人员流动）。本文预测制度化批判的结构更新率显著低于自发批判。若两者结构更新率相当，则「假性外部生产」被证伪——制度化的批判并不比自发批判更空转。

参考文献

- [1] Derrida, J. (2005). *Rogues: Two Essays on Reason* (P.-A. Brault & M. Naas, Trans.). Stanford: Stanford University Press.
- [2] Argyris, C. (1990). *Overcoming Organizational Defenses: Facilitating Organizational Learning*. Boston: Allyn and Bacon.
- [3] Merton, R. K. (1968). *Social Theory and Social Structure* (enlarged ed.). New York: Free Press.
- [4] Espeland, W. N., & Sauder, M. (2007). Rankings and reactivity: How public measures recreate social worlds. *American Journal of Sociology*, 113(1), 1-40.

关于本文的创新智商标注

本文在原始投稿基础上，由编辑依王德生《SDE本体论》与 SDE 创新智商评估规程做了「最近邻切割」与「可裁决预测」两节加固，意在抬高其不可还原性（I）与差异锐度（D）。依评分规程铁律，任何人不得为自己参与修改的文本认证分数，故本文所涉创新智商为**修改设计目标（134–138 区间）**，非认证分；认证分须由独立评分者盖出处另行给出。