

规则崩解：从“规则为交易而立”到“交易为规则而立”

AI经济底层构造原理的本体论翻转

作者 鲍锦朝 · SDE 学员 · 约 2.7 万字 · 发表于2026年7月24日

摘要 · ABSTRACT

互联网经济与AI经济之间的断裂，不应被理解为“规则形态从透明变为黑箱”，也不应被简单归约为“规则从可解释变为不可解释”。这些描述都还站在同一条河的同一边——它们都在问“规则长什么样”，却没有追问“规则本身是怎么被发生出来的”。本文提出，两者之间真正的断裂发生在更深的层面：规则的生成方式发生了本体论翻转。在互联网经济中，规则是“为交易而立”的——平台先制定规则，交易在规则划定的空间内发生。在AI经济中，这一逻辑被倒置：模型在无数次交易中自动抽取规律，交易本身成了规则的唯一源头。规则不再先行于交易，而是从交易数据中涌现。本文用一个新概念——“规则崩解”——来命名这一翻转的硬核：不是规则变得不透明了，而是“有一条独立的规则存在在那里”这个假设本身，在技术构造上失去了对应物。规则从未被“写定”在任何地方，它只是交易数据在高维参数空间中留下的一套痕迹。这套痕迹没有独立的身份，无法被抽取、被审计、被修订，因为它从未作为痕迹之外的什么东西存在过。这就引出了一个比“如何治理AI”更根本的问题：当“修改规则”这个动作失去操作对象时，那个被我们称为“治理者”的主体身份本身，它的存在条件正在发生什么变化？本文从规则的底座、交易的界面、市场的边界、主体的理性形态与产权-竞争结构五个维度，层层推证这一翻转如何从互联网经济自身矛盾的逼迫中发生出来，并正面辨析“痕迹”这一概念与复杂性理论中的“涌现模式”、德里达的“踪迹”以及STS学者关于“算法治理性”论述之间的不可通约之处。结论在给出核心论点的证伪条件之外，将追问推进到更深一层：在这场翻转中，连“需要治理”这个需求本身，都要被重新问题化。

关键词：关键词：规则崩解；算法治理；深度学习；在场权；痕迹

一、引言：当“修改规则”失去操作对象

在过去三十年的数字经济讨论中，有一个预设从未被正面宣告，却支撑了几乎全部关于平台治理、算法问责、数字监管的制度想象。这个预设是：经济活动的底层规则，无论多么复杂、无论由谁制定、无论以何种形态存在，它始终在那里——可以被定位、被调取、被审查、被修改。搜索引擎的排序算法有明确的权重因子，电商平台的信用体系有可追溯的评分条款，社交媒体的内容分发逻辑至少在原则上能被拆解为“如果-那么”的规则链。即使算法的复杂程度已超出单人即时理解，人们仍然相信——并且在监管话语中不断重申——这些规则“存在”于某处，以代码的形式被铭刻在服务器上，随时可以被调取、被审查、被修正。

这个预设同时承载着另一层更隐蔽的信念：规则在交易之先。不是时间上的先后，而是逻辑上的先后。平台必须先制定规则——什么商品可以上架、什么内容会被限流、什么行为触发封禁——然后交

易才在这些规则的框架内发生。规则是交易的立法者，交易是规则的臣民。这正是将互联网经济定性为“律令经济”的根据所在：平台的角色在结构意义上是一个立法者——它制定规则、执行规则、在规则被违反时施加制裁。这套运作逻辑与法律体系的同构性不是比喻，而是结构性的：规则是作为可被单独辨识、单独修改、单独审计的条文而存在的，就像法律体系中的法条。

本文不否认这一判断的洞察力。但本文要追问一个被它偷偷跳过的步骤：在互联网经济中，“有一条规则存在在那里”到底是什么意思？当人们说“搜索引擎的排序规则由数百个特征权重构成”时，规则是什么？是一段被写在某个配置文件里的代码。它可以是Python脚本中的一行“`if authority_score < 0.3: downrank(url)`”。它有一个物理位置（在某台服务器的某个文件夹里），有一个明确的身份（“第147行的那条降权规则”），有一部可以被追溯的修改史（这行代码上次是什么时候改的、谁改的、为什么改的）。规则是作为一件“东西”而存在的——可以被指向、被谈论、被争辩、被修订。它的存在方式是物件式的。就像一把椅子：你可以指着它说“这把椅子”，你可以把它搬走、修理、替换，它在所有这些操作中维持着它作为“这把椅子”的同一性。

这一存在方式，是律令经济一切制度构造的技术前提。因为规则是物件式的，所以它可以被审计——审计不是猜测系统“为什么这样做”，而是逐条比对“有没有这条规则”“这条规则是否被正确执行”。因为规则是物件式的，所以它可以被归责——当系统做出错误判断时，可以追溯到某条规则的缺陷、某个权重的错误设定。因为规则是物件式的，所以它可以被修改——“修改一条规则”这句话有一个明确的物理意义：找到那段代码，改写它，保存，重启系统。在这些操作中，规则始终维持着它作为“一条规则”的身份，就像一把椅子在被搬动时仍然是一把椅子。

这一整套制度想象，在AI经济的构造原理下开始从根部松动。本文要做的不是谨慎提醒这一松动的存在，而是正面论证：松动的原因不是AI“更复杂”或“更黑箱”——这些都是表层诊断，它们仍然在“规则存在在那里，只是我们暂时看不懂”的框架内运作——松动的原因更致命：在深度学习主导的经济规则生成中，“有一条规则存在在那里”这个句子本身，失去了它在互联网经济中的那个意思。那里没有“一把椅子”——那里只有一堆没有椅子形状的木头粉末、胶水和螺丝，它们以某种方式混在一起，形成了一个可供人坐下的表面。你可以坐上去，你可以说“这东西能坐”，但你不能指着任何一块粉末说“这把椅子”。规则从未以物件的形式存在过。

本文用一个新概念来命名这一断裂的硬核：“规则崩解”。它指的不是规则变复杂了或变模糊了，而是“规则作为一个独立存在物”这一存在论地位被取消。在AI经济中，规则不再是一个可以被单独指向、单独谈论、单独修改的物件。它崩解在模型的参数空间中——那里没有一亿个独立的“规则物件”，只有一亿个相互依赖的浮点数；没有一条一条可以被修订的条文，只有一次一次重新训练后整体漂移的权重分布。“修改规则”这个动作失去了操作对象——不是因为它更难，而是因为那个可以被捏住、被抽取、被改写的“它”压根就没有存在过。

本文的主线工作，就是给出这一崩解的精确发生机制，并论证：这不是一个暂时的技术局限，而是一次本体论翻转——经济规则的存在方式，从“物件式”翻转为“痕迹式”。物件先于使用而存在，痕迹在使用中才被留下。物件可以被修理，痕迹只能被重新踩踏。经济治理的全部制度工具——审计、归责、修订、解释——都建基于规则是物件；当规则变成了痕迹，这些工具就失去了它们的操作对象。但本文不打算在此处给出一个“治理药方”——那会把“崩解后的规则”重新当作一个发现学可以去诊断和给出疗法的现成对象。本文要做一件更难的事：追问那个被我们称为“治理者”的主体身

份本身，在痕迹化的世界里，它的存在条件正在发生什么变化。那个被认为要“去治理”的人——监管者、立法者、平台管理者——本身不也是被同一套发生过程配置出来的吗？当规则不再以物件的形态存在时，“治理”这个动作，作为一种社会存在，它还能从什么样的矛盾中被逼出来？

在展开这个论证之前，必须坦率交代这一判断的知识谱系。在更晚近的学术脉络中，Lessig (1999) 的“代码即法律”把代码视为一种特殊形态的规则——这仍然在“规则是物件”的范式内：代码是物件，可以被书写、被解读、被修改，它只是比法律条文执行得更自动化。本文的出发点恰恰在于：“代码是物件”这个预设，在深度学习那里失效了——不是因为深度学习不用代码（它当然用），而是因为深度学习产出的“规则”从未以代码条文的形式存在过。那些权重不是“被写定的”，而是“被算出的”。从“被写定”到“被算出”，规则的存在论身份变了。

同样，必须正面处理“规则崩解”与既有学术传统的关系。Aneesh (2009) 提出“algocracy”（算法治理）概念，已经敏锐地指出算法系统不通过“制定规则”来治理——这是他对科层制、市场与algocracy三种治理模式进行比较的核心论点之一。在algocracy中，治理不是通过成文的规则条文实现的，而是通过代码的系统性嵌入，在行为发生的层面直接调控可能性的空间。这一洞见与本文的出发点高度相关，必须坦率承认。但本文要追问的是：在algocracy的描述中，那套嵌入了代码、调控着行为空间的“东西”，它本身是以什么方式存在的？Aneesh讨论的是“算法如何执行治理功能”，但他未追问“执行治理的那套算法机制本身的存在论身份”。本文的核心推进在于：当深度学习主导经济规则生成时，那套调控机制不是“一套被写定的代码条文”，而是“一套从数据中自动形成的痕迹”——它不具备独立的、可被作为“条文”来指认的身份。这一辨别维度，是algocracy传统未曾切开的。algocracy说了“治理不再靠规则条文”，但没说“那套替代规则条文的东西，在存在论上是什么”。

与此同时，必须与另一脉同样切入“规则消解”议题的研究传统正面交锋。在过去十年间，以Mireille Hildebrandt和Antoinette Rouvroy为代表的STS学者，已经从不同路径触及了“算法治理正在消解传统法律规则”这一判断。Hildebrandt在其关于“算法审判”的著作中，系统讨论了基于数据的规范生成与传统法律规则之间的张力。Rouvroy提出的“算法治理性”（algorithmic governmentality）概念，则更为激进地指出，当代数据驱动的治理“不再通过规则，而是通过配置可能性的领域”来运作——它不颁布规范，而是实时调制行为发生的概率空间。这两个概念与本文的“规则崩解”共享同一个出发点：规则作为一条条独立的、可被写定的条文，正在丧失其在治理中的核心地位。但本文的辨别面在于：Hildebrandt和Rouvroy的论述，侧重的是“治理的操作逻辑”从规则走向了配置——她们描述的是治理“怎么做”发生了变化。本文追问的是“治理所依赖的那套机制，本身是什么”——它在存在论上是物件还是痕迹？换言之，本文的推进不在于发现“治理不再靠规则条文”（这一点在algocracy和算法治理性文献中已被充分论述），而在于揭示：那套替代规则条文、承担治理功能的机制，在本体论上已经丧失了一切可以被称作“条文”的存在方式。它不是“更隐蔽的条文”，它不是“更难解读的规范”——它根本就不是任何一种规范。它是痕迹：是从数据中残留、从未被写定、也无法被独立调取和修订的参数分布。这一区分将“算法治理”的讨论从对治理技术的功能描述，推进到了对治理技术自身存在方式的追问。Rouvroy说治理“不再通过规则”，本文说的是：治理当前所依赖的那套东西，本身连“不是规则的规则”都不算——它不具备任何规范物的身份。

此外，复杂性理论中的“涌现”概念也需要被正面回应。涌现论早已指出，规则可以从大量局部交互中自下而上地生成，而不需要自上而下的制定。本文的“规则从交易数据中涌现”，很容易被读作涌现论的一个应用案例。但本文的辨别面在于：涌现论即使强调规则的自发生成，在多数论述中仍保留着一个隐含预设——涌现形成之后，规则可以作为某种“模式”被辨识、被研究、被对象化。圣塔菲研究所讨论的“涌现秩序”，仍然被视为一个可以被建模、被描述、在一定程度上被调控的对象。正是在这一点上，本文的“痕迹”概念切开了与“涌现模式”之间的一道硬缝。涌现在稳定之后的产物，具有一种“可被辨识的形态”——你可以指着沙堆上反复出现的某类崩塌模式说“这就是 $1/f$ 噪声”，然后把它的统计特征写入一篇物理学论文的公式里。模式具有客体身份：它有名称，有可被量化的特征，有可被其他研究者独立复现的可观测性。痕迹则不同——痕迹即使在涌现稳定之后，仍然不具备这种客体身份。深度学习模型参数不是沙堆崩塌的模式。沙堆崩塌的模式你可以用一台摄像机拍下来，圈出来，放进教科书里。深度学习模型的行为你也可以录像——但那个让你想去“拍下来”的东西，不在录像里，不在任何一行代码里，也不在任何一组可以圈出来的参数里。它是残留，不是轮廓。它没有可以被指认的边界。你可以感受一个推荐系统的“整体倾向”，但你不能指着任何一组参数说“就是这个倾向”。痕迹是彻底的不可对象化——这就是本文对涌现论的关键推进：不是在涌现的自发生成层面与之一致，而是在涌现产物的存在论身份上与之一刀切开。涌现论走到“自发生成的模式”就停了；本文把追问推进到那一步之后，发现那套东西连“模式”都算不上。

可解释AI (XAI) 文献中关于LIME、SHAP等事后解释工具的讨论，本文也将正面回应：这些工具不是在“翻译”模型的内部规则，而是在模型表面贴一层物件式的薄膜——薄膜下面的痕迹本身，依然不是物件。

二、规则的底座：从物件到痕迹

2.1 物件式的规则：互联网经济的运作前提

要理解“物件式”规则的运作逻辑，最好的切入点是看一个互联网平台如何设计它的搜索引擎排序规则。在典型的搜索引擎架构中，排序逻辑由数百个特征构成：页面的关键词匹配度、网站的权威性评分、内容的时效性、用户点击率的历史表现、页面加载速度等等。每一个特征都被赋予一个明确的权重。关键不在于这些权重是简单还是复杂——关键在于它们作为物件的存在方式。权威性评分的权重可能是0.35，这个数字被写在一个配置文件中，通常是一个YAML或JSON格式的文件，存储在平台代码仓库的某个目录下。一个产品经理可以打开这个文件，找到“authority_weight”这个字段，把0.35改成0.28，保存，提交一个代码改动申请，经过审批后合并到主分支，然后部署上线。在这一整套操作中，“0.35”这个权重始终是一个可以被指向、被操作、被追踪的物件。它的修改史被记录在代码仓库的提交日志里——2023年6月，张三修改了权威性权重，从0.4降到0.35，因为用户反馈显示高权威网站淹没了时效性内容。

这种物件式的存在方式，是律令经济的核心承诺得以兑现的技术基础。可审计性——因为规则是物件，所以审计师可以打开配置文件，逐条比对：这条规则写的是什么？它的执行逻辑是什么？它在什么时候被触发？可归责性——因为每条规则都有明确的“主人”（制定它的团队、修改它的工程师），当规则出错时，可以追溯到具体的人、具体的改动。可修订性——“修订”这个动作有一个明确的物理意义：改动那个物件，保存，部署。平台治理的全部制度工具——从内容审核指南到广告政策到数据隐私条款——都是建基于规则的物件性。规则可以是错的，规则可以被滥用，规则可以变得过于复杂以至于难以被整体把握——但它至少在那里。治理的全部努力，都在于让这件东西更透明、更公正、更有效。没有人质疑过这东西本身的存在。

2.2 物件范式的极限：被自身逻辑逼到临界点

物件范式的失效，不是在某一个戏剧性的时刻被“废弃”的。它是被自身的逻辑逼到极限后自动坍塌的。在平台的初创期，规则可能有几十条。每条规则都有一个清晰的意图，规则之间的交互效应可以被少数几个资深工程师勉强把握。但随着平台规模的增长，规则从几十条变成几百条、几千条。到这个阶段，“逐条修改”这个动作的意义开始发生质变。修改一条信用评分的规则，可能间接改变了反作弊规则的触发阈值，进而影响了内容推荐算法对“优质内容”的判断——而这三个规则是由三个不同的团队在不同时间写下的，没有任何一个人能完整推演出这一连串的连锁反应。当一个系统内有三千条互相触发的规则时，“修改第147条规则”这个动作，在物理上仍然可以被执行——你可以找到那行代码，改掉它——但在这个动作与其实际效果之间，已经不存在一条可被推演的因果链。你改了那行代码，你不知道会发生什么。

这时的规则，在名义上还是物件——它们仍然以代码条文的形式存在于配置文件中——但在运作上已经不再是物件了。因为“作为一个物件”的核心特征，就是它可以被单独操作而不丧失其同一性。一把椅子可以单独被搬走，而房间的其余部分不受影响。但当三千条规则相互缠绕到“动一条则全局漂

移”的程度时，规则就不再可以被单独操作。它们虽然在物理上还是分开写着的，但在功能上已经长成了一个纠缠的整体。

这就是物件范式被自身逻辑逼到临界点的具体机制。律令经济预设“规则是物件”——可以被单独制定、单独修改、单独审计。但当规则的交互效应超过人类可逻辑推演的上限时，“单独修改”就从一个真实的操作变成了一种自我欺骗——你以为你在改一条规则，其实你在扰动一个整体。物件性的外表还维持着，但物件性的实质已经萎缩了。这一萎缩的根本驱动力，不只是“规则变多了”这个数量问题，而是随着规则数量的增长，维持物件式规则的边际成本呈指数级上升。每增加一条规则，不仅要调试它本身的逻辑，还要测试它与所有既有规则之间的交互效应。当规则数量突破某个临界点时，维护一个物件式规则体系的人力成本和出错成本，就超过了经济上可承受的上限。这不是一个“更好的规则管理工具”可以解决的问题——因为交互效应的数量与规则数量的平方成正比，它本质上是一个不可被线性管理工具征服的复杂度爆炸。在这个临界点上，端到端训练的深度学习不是作为一个“更好的选项”出现的，而是作为在此复杂度量级下唯一幸存的技术路径出现的——它放弃了对每一条规则的独立掌控，以此换取了在整体性能上的持续提升。这不是一个有意的选择，而是结构性逼迫的结果：维持不住物件身份了，就只能松手。

需要特别强调：这个临界点不是AI制造出来的。它在深度学习大规模进入经济系统之前，就已经被互联网经济自身的规模扩张逼到了眼前。AI的介入不是制造了这个问题，而是把律令经济已经在自身极限处被迫经历的这一次坍塌，推进到了它的逻辑终点：连“物件”这个外表都放弃了。在新的构造原理下，规则从未以物件的形式存在过——它们从一开始就是痕迹式的。

2.3 从物件到痕迹的一次具体翻转：RankBrain如何改变了“修改规则”的含义

概念演绎无法替代技术史追踪。本节以谷歌2015年将RankBrain引入核心搜索排序算法这一真实历史技术事件为案例，逐节追踪“规则从物件翻转为痕迹”是如何在一次具体的工程决策中发生的。

在RankBrain被引入之前，谷歌的搜索排序算法由数百条物件式规则构成。PageRank、关键词匹配度、页面新鲜度、站点权威性——每一条都对应着一组可被写定的权重。工程师可以打开配置文件，找到“freshness_weight”这个字段，把它从0.25调成0.3，然后部署。这个修改动作的含义是清晰的：新鲜度在整体排序中的重要性被调高了。如果出错，可以回滚。

2015年，困境出现了。谷歌面临一个具体的工程压力：每天有约15%的搜索查询是谷歌此前从未见过的（Google, 2016）。这些长尾查询——通常是多个词组成的自然语言问句——无法被基于关键词匹配的物件式规则有效处理。传统做法是：当出现一个新查询时，工程师分析查询意图，手工编写一组关键词匹配规则。但这个做法在面对每天数以亿计的新查询时彻底失效了——没有任何团队能以这种速度编写规则。

困境逼迫了一次架构决策。谷歌工程师决定引入一个基于深度学习的查询理解系统——RankBrain。RankBrain不做关键词匹配。它将整个查询编码为一个高维向量——一个有数百个维度的浮点数数组。然后将这个向量与排序算法中的其他信号共同输入最终的排序公式。这里有一个关键的技术细节：RankBrain输出的向量不是作为一个“独立规则”存在的——它不与任何一行配置文件的任何一行“if-then”语句对应。它是在模型训练过程中，从海量查询-点击数据中自动学出来的。修改

RankBrain的“规则”的唯一方式是重新训练整个模型——用新的数据跑一轮新的训练，旧的权重分布被整体替换。

在这个时刻，“修改一条排序规则”这个动作，对于通过RankBrain处理的那些查询来说，失去了它在物件式规则时代的那种含义。你不能说“把RankBrain对长尾查询的敏感度调高三档”——因为那里没有“一档”。模型的参数是连续分布的高维向量，不存在可以被独立调取的“长尾查询敏感度”这个控制杆。

这一翻转逼迫工程师团队面对一件在物件范式下可以轻松做到、在痕迹范式下突然做不到的事：当某个广告主投诉“为什么我投了广告，我的页面在这个查询下还是排在第三页”时，工程师无法给出一个物件式的解释。他们不能打开配置文件，找到某行代码，然后说“是这条规则导致你的页面没有被排上去”。他们能做的只是分析：RankBrain为这个查询编码出的向量，与这个广告主页面的内容向量之间，在各种维度上的“距离”是多少。但这个“距离”不是一个可以被叫作“规则”的条文——它是一亿次矩阵乘法的联合作用的结果。它不是被写定的，它是被算出的。

RankBrain的意义不在于它“更聪明”。它真正的破坏性在于：它在谷歌核心搜索产品的核心位置，让“规则作为一件可以被单独指向的物件”这个预设，第一次在工程上失去了对应物。此前搜索引擎的架构是一套“物件式规则体系”加上“一些处理边界情况的启发式方法”；此后，它变成了一套“痕迹式参数分布”加上“一些包裹在痕迹外层的物件式规则”。这个翻转不是理论的推演，而是2015年到2016年间在谷歌办公楼里实际发生的工程实践。

2.4 痕迹式的规则：深度学习如何生成经济秩序

RankBrain案例展示了翻转如何发生。本节从技术原理层面，解释为什么深度学习必然产出痕迹式的规则——这不是选择，而是深度学习达到当前性能所依赖的必要条件。

深度学习模型的核心技术构造——分布式表征与端到端训练——在结构上排斥“独立可操作规则模块”的设计可能性。分布式表征意味着：模型习得的任何一个“概念”（如“权威性”），不是被存储在一个特定的参数或一组特定的参数中，而是以高度分散的方式分布在全部参数的取值模式里。你无法把“权威性识别”从模型的其他功能中“拆卸”出来单独修改，因为没有任何一组参数单独对应“权威性识别”。端到端训练意味着：模型的所有参数是同时被调整的——调整的目标不是让某一组参数变得“更准确”，而是让整个模型的输出与人类标注之间的整体误差最小化。在这个过程中，参数之间形成了极其复杂的相互依赖关系，任何单个参数的值只有在与其他所有参数的联合作用中才有意义。一个推荐模型的权重没有独立性。你不能把第1,478,332个参数单独拿出来，说“这个参数的含义是什么”——因为它只有在与其他一亿个参数的联合作用中才产生效果，单独抽出来，它就是一个没有意义的浮点数。

这就是“痕迹”这个隐喻的精确含义。你在沙滩上走过，留下了一串脚印。这串脚印承载着你的体重、步幅、行走方向的信息——如果你愿意，你可以从脚印里“读出”这些信息。但脚印不是物件。它不是在沙滩上事先放好的形状；它纯粹是你走过之后的残留。你没有在任何一刻“写下”脚印——你只是走过，然后脚印就留在了那里。深度学习模型的权重，就是这类痕迹。它们是模型在海量数据上“走过”——在数值优化算法的一轮又一轮迭代中，不断调整自身以逼近最优预测——之后留在参数空间里的一套残留。这套残留承载着数据中的统计规律，但它从未被任何人“写定”为规则条

文。“模型习得”不是“某个人制定了规则”，甚至不是“某个人设计了规则生成的流程”——那种流程本身就预设了规则最终以物件的形态出现。“模型习得”是：让算法在海量数据上跑，跑完以后，参数空间里自然留下了一套痕迹，这套痕迹能指导模型的行为。规则从未被制造出来，规则只是被留下了。

这就是物件与痕迹的**本体论差异**。物件是一个独立存在物——它的存在不依赖于任何具体的使用情境，你可以把它从环境中抽出来单独考察、单独修改、单独替换。痕迹则永远附着在它被留下的那个介质上——你不能“单独”抽出一枚脚印去研究，你必须在整片沙滩的语境中去看它，因为它与旁边的脚印、沙的湿度、踩踏的力度共同构成一个不可分割的整体。这就是“规则崩解”的硬核：规则不再是一个一个的物件，可以分别被指向、被谈论、被操作。规则崩解成了一片痕迹——整体有效用，个体无身份。

2.5 痕迹不是什么：与“涌现模式”和“踪迹”的正面辨析

走到这里，“痕迹”这个概念已经暴露在一个危险的位置上。它很容易被读作复杂性理论中“涌现模式”的别名——一种从局部交互中自发生成、整体不可还原的稳定构型。它也很容易被联想为德里达笔下那个同样反物件、反现成性的“踪迹”（trace）。如果本文只是用一个新的术语去指涉一个已经被充分论述过的现象，那么“规则崩解”的理论收益将大打折扣。本节必须正面完成这个辨析：指出“痕迹”在哪一个具体的辨别面上，切开了“涌现模式”和“德里达的踪迹”所没有切开的维度。

首先看涌现模式。以圣塔菲研究所为代表的复杂性科学传统，早已反复论证过一条原理：从鸟群的编队飞行到沙堆的崩塌分布，稳定的宏观模式可以从大量局部交互中自下而上地涌现，而不需要任何中心化的规则制定者。涌现出的模式一旦稳定，它就具有“可被辨识的形态”——你可以拍下鸟群在空中的编队形状，圈出它，说“这是V字形编队”；你可以录下沙堆崩塌的声学信号，分析其中1/f噪声的频谱特征，写进教科书。模式具有客体身份：它有名称，有可被量化的特征，有可被其他研究者独立复现的可观测性。你可以指着它的某些属性说“这里”——那里有一个“沙堆的临界角度”，它不是一个物件，但它是一个可以被精确测量和独立谈论的量。

深度学习模型的参数空间完全不是这样的。没有人能指着那个一亿维的参数空间中的任何一维、任何一组维度，说“这里”——因为没有一个参数或一组参数单独、稳定地对应任何一个可被辨识的“模式”。端到端训练的代价正在于此：它放弃了把任何“模式”单独凝结成可被指认的客体的可能性。沙堆上反复出现的崩塌模式，会把沙粒按某种统计规律重新排列，但你可以圈出那个崩塌的模式说“它每次大概长这样”。深度学习模型的“规则”则不同——你没有办法圈出任何一组参数说“这一块管着权威性识别”。不是说它更难圈，而是“圈出”这个动作在这里没有任何对应物。这不同于涌现模式：涌现模式一旦涌现，它就可以被作为模式来对象化——你可以用统计量描述它，可以用模型复现它，可以在它偏离预期时探测到偏离。“痕迹”则不同：痕迹没有可以被对象化的稳定轮廓。它确实是某种残留，但它从来没有凝结成一个可以被独立指认的“形状”。

再看德里达的“踪迹”。德里达终其一生都在用“踪迹”这个反概念去松动西方形而上学中“现成在手的存在”这一预设。踪迹不是物，不是概念，不是在场，总已经先于一切在场而在与不在场的交织中被涂抹、被擦除。本文使用“痕迹”一词，确实与德里达的踪迹有家族相似：两者都拒斥物件式的独立存在方式，都强调其不能从生成过程中被抽出。但有一个决定性的差别。德里达的踪迹是“总已

经在”的——它不是被任何具体过程“留下”的，它是一种比一切具体生成更原初的差延运作。踪迹没有一个可以被经验地追溯的“被留下的过程”。本文的“痕迹”恰恰相反：它的存在论身份，全在于它是“被留下的”。它有一个可以被经验地追溯的发生过程——一次具体的训练、一套具体的数据、一个具体的优化算法。这不是某种先于一切存在者的原初涂抹；这就是2023年某一天，某个数据中心的某个GPU集群在跑了若干小时后，留在参数空间里的一套具体的残留。它具有绝对的经验性——你可以在技术上复现它、覆盖它、比较两次训练留下的不同痕迹。它不是任何先验的差延结构，它就是深度学习这个特殊的工程实践所产生的一类具体的、可经验地追踪的存在物。

正是这个差别，决定了本文的“痕迹”不能被德里达的“踪迹”所替代。踪迹的任务是松动形而上学的根基；痕迹的任务是揭示一种特定的技术实践——端到端训练的深度学习——正在制造出一类无法被既有治理范畴收束的新型存在物。两者的哲学同盟在于都反物件；但两者的辨别面在于，痕迹有一个可以被具体技术史追溯的发生条件，而踪迹没有。这是本文在哲学概念库中为“痕迹”划定的独有坐标。

2.6 痕迹无法被“翻译”回物件：可解释AI的本体论边界

可解释AI（XAI）作为一个活跃的研究领域，其隐含的承诺是把痕迹“翻译”回物件——用事后解释工具（如LIME、SHAP、显著性图）为模型的输出生成一套人类可理解的“规则说明”。如果这个承诺能兑现，那么“规则崩解”就只是一个暂时的技术状态——等解释技术足够好，痕迹就可以被重新还原回物件。

这个承诺在经验层面已经遭遇了系统性的困难。当前最先进的事后解释工具能够在特定输入-输出对上给出事后的近似解释，但不同的解释工具对同一个模型的同一个决策，经常给出相互矛盾的归因。例如，在对一个信用评估模型拒绝某笔贷款的决策进行解释时，基于SHAP值的解释可能会指出“收入水平”是主要的负面因子，而基于LIME的局部近似模型则可能指出“职业类别”才是决定性因素。这不是技术尚未成熟的暂时现象，而是结构性的——这些解释工具不是在“翻译”模型的内部决策逻辑，而是在用一个简单模型去“拟合”复杂模型在某个局部的行为。更关键的证据来自医疗AI领域。已有研究者公开记录了深度学习医学影像诊断系统中反复出现的一类现象：当研究人员尝试根据显著性图的反馈去“修正”模型时——以为模型在“错误地关注了不相关的区域”——模型的诊断准确率反而下降了（参见Ribeiro, Singh & Guestrin, 2016对LIME局限性的原始讨论；Lipton, 2018对“可解释性”概念本身的批判性检视）。这说明显著性图所揭示的东西，并不等同于模型实际做出诊断时所依赖的东西。解释工具不是在“翻译”痕迹——它们是在痕迹的表面贴物件式的标签。

为什么这种翻译在本体论上是不可能的？因为“规则”这个词在物件范式有一个预设：规则是被制造出来的。它是一种行为指令，预先设定好了“在什么条件下做什么事”。物件式的规则有一个指令身份——它是规范性的，它告诉系统“应该怎么做”。但深度学习模型的权重不是规范，而是残留。它们没有指令身份——它们不告诉系统“应该怎么做”，它们只是系统在训练过程中被调整成那样的状态。痕迹是描述性的——它描述的是“系统在数据上学到了什么规律”，而不是规范性的“系统被命令做什么”。你可以在事后给痕迹贴上一个规范性的标签——“模型似乎学会了‘权威性低的网站应该降权’”——但这个标签是你赋予的，不是痕迹自身携带的。痕迹自身不携带任何“应该”。它只是在那里。

这就是痕迹永远无法被无损翻译回物件的根本原因。翻译预设了两边指涉的是同一种东西——就像把“law”翻译成“法律”，因为两边都是物件式的规则。但当你把痕迹“翻译”成物件时，你做的不是转换语言，而是强行给一件没有指令身份的东西赋予一个指令身份。这就像你在沙滩的脚印旁插了一块牌子，写“此处规定行人应走此方向”。那串脚印不是规定，它是走过之后留下的。它承载着过去的行为信息，但它不规定未来的行为。同样，模型权重的修改方式——“重新训练”——是痕迹式的：你不可能“修订”一串脚印，你只能在沙滩上重新走过一遍。旧脚印会被新脚印覆盖。没有修改史，只有覆盖史。“一条规则是如何演化的”这个问题，在一个物件式规则的世界里是合法且有答案的；在一个痕迹式世界里，这个问题本身就是对范畴的错用——那里没有“一条规则”可以演化，只有整片参数空间的漂移可以被观察。但这个“被留下”的过程本身，在不同的技术架构与商业压力下，其具体形态又可能发生什么样的变异？在不同的训练数据分布、不同的模型架构选择下，那套痕迹的“黏稠度”与“顽固度”是否会有所不同？把这个追问钉在这里，是为了不让“痕迹”这个词变成一个静态的终局命名——它的形态本身，还在不断地被新的矛盾逼着继续形变。

三、交易的界面：绵延交互中，不再有“一笔交易”

3.1 匹配器的逻辑：交易作为可被清点的物件

在互联网经济的规则体系下，不仅规则是物件——交易也是。平台的核心功能是匹配：将供给方与需求方撮合到一次交易中。这个“一次交易”不是自然发生的经济行为，而是代码在连续的行为流中强制制造的一个断点。用户在电商平台上的购买行为在物理世界中是一个绵延的过程——她可能几天前看到一条广告，昨天跟朋友聊天时聊到这个品类，今天上午用手机搜索了一会儿，午饭后又打开APP继续浏览，最后在下午三点下单。在这个绵延中，代码打下了一系列时间戳：“用户点击了广告”——一个事件被记录；“用户搜索了‘蓝牙降噪耳机’”——又一个事件；“用户提交了订单”——交易达成。每个事件在服务器日志中占一条记录，每条记录都有明确的时间、明确的动作类型、明确的主体身份。这些日志是平台经济学的原子——不可再分、边界分明、可被清点。

律令经济之所以能建立一套关于交易的经济学——点击率、转化率、客单价、归因链条——正是因为代码把绵延切割成了离散的事件物件。每一个事件物件都可以被独立地统计、分析、归因、征税。就像规则物件使得审计成为可能一样，交易物件使得经济度量成为可能。经济学家的全部统计工具，都建基于一个本体论预设：经济活动可以被不失真地切割成一个个独立的交易单元。如果一个经济体中的“一笔交易”无法被客观地指认出来，那么这个经济体的宏观经济数据就失去了解释对象——不是数据不准，而是在那里根本没有一个叫做“一笔交易”的东西可以被度量。

3.2 切割的任意性：多轮对话如何暴露匹配器的预设

生成式AI并没有“推翻”匹配器，而是让匹配器所依赖的那种切割动作，在对话交互中暴露出了它的预设性。匹配器预设：每一次交易都有一个可以被客观确定的起点和终点。但在多轮对话中，这个预设无法被兑现。

考虑一个基于已有公开案例与行业典型交互逻辑进行综合重构后的分析性场景。[1] 在这一场景中，一个主流电商平台在其移动端部署了基于大语言模型的AI导购系统，以下为其中一次具有代表性的交互还原。用户打开对话界面，输入的第一句话是：“我想给父母买一个操作简单的血压计，但是他们不太会用智能手机。”这不是一个可以被任何传统搜索引擎精确匹配的查询词——它不是问“什么品牌的血压计好”，也不是问“血压计多少钱”，它是在描述一个需求的情境。AI助手的第一轮回复是询问预算范围和父母的大致年龄；第二轮推荐了两个不同型号并解释各自的操作难度；第三轮，用户犹豫时，AI主动追问“您父母平时是自己测还是社区医生上门测”——这个追问的意义不在于获取更多信息，而在于它暗示AI正在认真考虑使用场景的细节；第四轮，用户忽然岔开话题问“对了，最近是不是有新的高血压用药指南出来”——AI接住了这个岔开的话题，提供了相关信息，然后在某个自然的话轮把话题带回血压计选购；第五轮，AI提供了一段详细的对比分析，用户说“那就第一个吧，帮我下单”。

在这个完整的交互序列中，哪一个时间点是一笔“交易”的起点？是第一句话“想给父母买血压计”吗——但那时用户甚至没有明确要购买，这些话也可能随后转向“其实也不急，我先了解一下再说”。是第五轮用户说“帮我下单”吗——但那句话是整个交互过程的自然收束，它的意义是由前面

全部话轮共同赋予的。把“帮我下单”抽出来孤立地标记为“购买事件”，就像把一段旋律的最后一个音符标记为“旋律”——你不是在记录一个自然存在的单元，而是在连续体中强行制造一个断点。更关键的是，用户对AI的信任在第三轮发生了一次微妙的变化——那个追问“谁测”的动作，在用户感知中可能被解读为一种体贴的信号，而这个被感知到的“体贴”在第四轮用药指南的旁涉中进一步加深。到了第五轮，当AI给出对比分析时，用户说“那就第一个吧”的那个“就”字里，已经沉淀了前四轮累积的全部信任。这个信任的生成过程不能被拆成独立的步骤——它是绵延的、相互渗透的、在每一轮对话中都同时被前一轮塑造并塑造着下一轮。（需要明确：本节对这一交互中“信任”变化的判断，是基于对话逻辑的叙事重构，而非基于情感计算或经验实验的断言。它服务于呈现论证逻辑，不声称其经验精度。）

匹配器的切割动作，在这里暴露了它的预设性。它预设“交易”是一个自然存在的事件单元，代码只是忠实地记录它。但在这个对话中，没有任何一个点可以被客观地指认为“交易发生之处”。“帮我下单”那句话之所以是交易，仅仅是因为代码选择在那里打下了一个时间戳。代码不是在记录交易，而是在创造交易——它强行把一段绵延切成了“对话部分”和“交易部分”，然后把后者标记为“一笔交易”。这个切割是任意的，它不是对客观存在的交易原子的发现，而是对绵延的暴力截断。

从这个矛盾中，交互的本质被逼着浮现出来：交易不再是以物件的形态存在，而是以绵延的、弥漫在整段对话中的方式存在。它不发生在某一个点上，它撒在每一轮对话里。匹配器之所以在这种交互中失效，不是因为匹配器不够好，而是因为匹配器要切割的东西——“交易物件”——在这种交互中根本不存在。在这个意义上，交互本身构成了一个绵延的场。

3.3 当“一笔交易”不再是物件：“监管”这个身份发生了什么

当交易从离散物件蜕变为绵延的交互场，产生了一个比“监管如何适应”更根本的问题。律令经济中的监管，本质上是物件的清点工作。反垄断监管要检查平台是否对自己的商品给予了不公平的曝光优势——方法之一就是比对不同商品在搜索结果中的展示位置，因为每一次展示都是一条可以被清点的日志。所有这些监管技术共同依赖一个前提：存在可以被单独提取、单独审查、单独判断真伪的“行为物件”。在绵延交互中，这个前提不再天然成立。

我们可以从一次真实的算法审计实践来看这个困境的具体形态。ProPublica对COMPAS累犯预测算法的调查遭遇的核心困难，在本体论上与本文讨论的困境同构：COMPAS的预测不是基于一条条可以被单独调取的物件式规则——它基于一套从历史数据中训练出来的统计模型。审计团队无法打开这个模型，找到“种族”这个特征，看它被赋予了多大的权重，因为模型内部没有这么一个单独的权重。他们只能用外部统计方法——比较不同种族被告人的预测准确率差异——来间接推断算法是否“考虑了”种族。但即使发现了差异，他们也说不清这个差异来自模型的哪个参数、哪条规则、哪次训练决策。他们能做的只是对模型的行为做事后统计，而不是对模型的内部规则做审计。

这就是“清点物件”的监管逻辑撞上“只有痕迹、没有物件”的AI系统时的真实处境：监管者想要检查的是一个物件，但那里没有物件。他们只能绕着系统的外围打转，测量输入与输出之间的统计关联，却永远无法打开机器的外壳，找到那条写着“if race == ‘black’ : increase risk score”的代码行——不是因为那条代码藏得太深，而是因为它从未被写过。

这里有一个常见的诱惑：把这个困境表述为一种“监管失灵”，然后自然地转向对“新监管工具”的呼吁。本文必须抵抗这个诱惑。发生学要追问的不是“监管者面对这个新现象该怎么办”——那个问题已经把“监管者”当成一个站在现象对面的、不变的主体了。本文要问的是：在这个困境中，被重新配置的不只是“监管的对象”，连“监管者”这个主体身份本身，它的存在条件也在发生变化。当交易不再以物件形态存在时，那个在物件形态所支撑的制度框架中被建构出来的“监管者”——它的专业知识、它的管辖权边界、它的审计工具、它与被监管者之间的可辨识关系——这些构成“监管者”身份的存在条件，还能在原地维持吗？如果不能，那么“监管”这个动作，作为一种社会存在，它将从什么样的新矛盾中重新被逼出来？把这个追问推进到这里，不是为了给出答案，而是为了指明：在痕迹式的世界里，“治理”不是一个先在的、只是在等待新工具的主体——它本身也是被发生的。

四、市场的边界：当二值逻辑失去可操作性

4.1 硬性的边界：二值逻辑与物件式划分

互联网经济的市场边界，具有一种类似于围墙的实体性。这个边界不是物理的，但同样是物件的：电商平台明确规定哪些类别的商品可以上架、哪些不可以——武器不可以，酒类需要特许资质，虚拟商品有特殊的交付规则。这些边界之所以是“硬”的，是因为它们被写在代码里，并且代码的执行具有二值逻辑：一条内容要么被标记为违规（并因此被自动下架），要么没有；一个商品要么在可售类目标清单里，要么不在。边界的存在是二值的——商品要么在墙内、要么在墙外。这种二值性之所以能够运作，是因为边界本身是作为物件被写定的。平台制定了一条规则——“禁止售卖武器”——这条规则是物件式的：它有一个明确的语义，可以被翻译成代码中的条件判断语句（“if category == ‘weapon’ : reject()”），在每一次商品上架时被机械执行。

4.2 二值逻辑在自我侵蚀中被逼出“弹性分层”

二值逻辑的困境不是AI“制造”出来的——软化的边界并不是因为AI更聪明、更灵活，而是二值逻辑在自身的规模化运作中必然积累出其无法消化的矛盾后果，从而逼迫一种新的边界形态从矛盾中发生出来。

当一个深度学习模型被用于评估内容风险时——例如判断一条用户发言是否构成“隐性骚扰”——这个判断不再是二值的。模型输出的不是一个“是/否”标签，而是一个概率分布：这条发言有72%的概率构成隐性骚扰。在硬边界的逻辑下，这个72%必须被翻译成一个二元判决：杀，还是不杀。杀，意味着28%的概率是在误杀；不杀，意味着72%的概率是在漏放。无论选哪一个，硬边界在概率性的灰色地带做出的二元判决都会积累矛盾后果——误杀率与漏放率在成千上万次决策中持续攀升，用户申诉量爆炸，审核团队被淹没，信任基础在反复的误杀与漏放中被持续侵蚀。

这逼迫出一种变化。在2019年到2021年间，多个主流社交平台的内容透明度报告记录了用户申诉量的急剧攀升。在这个压力下，审核团队开始自行发明出一种中间状态：对判定为“高概率但不够确定”的内容，不打“违规”标签，而是打“待观察”，推迟处置，等待更多上下文信号。这个中间状态的发明，最初是一个操作惯例，不是一个被产品经理设计出来的正式功能。它是由审核员在手忙脚乱中摸着做出来的。但正是这个从操作层被逼出来的惯例，逐步固化成了系统架构的新常态：模型输出从“违规标签”变成了“风险分数”，内容分发不再依赖二值的“放行/拦截”开关，而是依赖一条渐变的“风险分层”光谱；在此基础上，系统又进一步引入了“用户信誉分层”——信用高的用户发布的内容走快速通道，信用低的走深度审核。至此，那个最初只是在操作层被逼出来的“中间状态”，在系统架构中被固化。边界不再是二值的硬墙，而是渐变的风险分层。它不是任何一个人“设计”出来的——它是在二值逻辑的刚性面对概率化决策的规模压力时，被逼着长出来的一层弹性响应。

4.3 从硬墙到弹性分层：边界的痕迹化

这一过程揭示了边界痕迹化的发生机制。在物件范式下，边界是一件东西——它被制定、被摆放、被修改，在这些操作中它维持着作为“边界”的同一性。你可以指着电商平台的违禁品清单说“这就是市场边界”。但在AI经济中，边界是在二值逻辑的自我侵蚀中被逼出来的新形态——它不是任何一个人有意“铺设”的，而是模型、数据、用户行为、审核团队的操作惯例在持续交互中共同“养”出来的一层弹性分层。

要理解这一点，需要回到AI模型做决策的方式。假设一个搭载大模型的智能助手被用于筛选和推荐理财产品。监管规则规定，某些高风险产品只能推荐给“具有一定风险承受能力的合格投资者”。在物件范式的架构中，这条规则会被翻译成一系列硬性指标：年收入不低于X、金融资产不低于Y。边界是一条清晰的线。但一个深度学习模型处理同样的问题时，它的判断逻辑完全不同。它从用户对话的语言风格、浏览历史的整体模式、在平台上的交互节奏中，捕捉到的不是资产数字，而是一个复杂的概率分布：这个用户在实际投资行为中表现出理性判断力的概率有多高？如果概率高于某个阈值，推荐；如果低于，则不推荐。这个阈值本身也不是硬性设定的——它在模型训练中被自动调整，并且在持续的在线学习中可能继续漂移。

边界不再是一条“年收入正好等于X”的线，而是一个连续变化的概率值。它没有一条可以被指认为“边界规则”的条文——它是模型在训练过程中自动形成的一种隐性倾向。边界是流动的——随着模型的持续训练，整个市场的风险分层在缓慢地漂移，这周还能被推荐某产品的用户，下周可能就不再被推荐，不是因为用户的资产状况变了，而是因为模型在新一轮训练中调整了它对风险信号的敏感度。这就是边界的痕迹化：边界不是被写定的物件，而是被踩踏出的痕迹。你无法单独调取“边界规则”来修改它——你能做的只有重新训练模型，让新数据在参数空间里踩出新的一层痕迹。

五、从“记录权”到“在场权”：产权逻辑的翻转

5.1 记录权：归因链作为产权的基础

在互联网经济中，产权归根结底是“记录权”。这句话的意思是：谁拥有记录一条经济行为之因果链的能力，谁就实质上拥有对这一行为之经济价值的解释权与索取权。这种记录权的技术载体是日志。每一次点击、每一条搜索、每一笔交易，都在平台的服务器上留下一条或多条结构化的日志记录。广告平台之所以能向广告主收费，不是因为它在概念上“帮助品牌获得了曝光”，而是因为它可以递给广告主一份清晰的归因报告：用户A在周三下午3:12分点击了你的广告，然后在周五上午10:38分下单了你的产品——这两条日志之间的因果链，被平台的归因模型串联起来，成为广告费结算的依据。记录权之所以能在互联网经济中运作，是因为它有两个技术前提。第一，交易是物件——每一次交易都被代码切割为一个独立的、可以被记录的事件。第二，归因链是物件——两条日志之间“点击→转化”的因果关系，被归因模型以明确的规则（如“最后点击归因”）固定下来，成为一条可以被审计的记录。

5.2 在场权的发生：从一次归因失败的逼迫中长出

记录权在AI经济中面临一个结构性的困境。当交易的界面从离散的匹配事件蜕变为连续绵延的交互场时，日志所记录的事件序列，就不再能串起一条清晰的、可以被单一主体所独占解释的因果链。

回到本文第三部分所描述的那个AI导购购买场景。用户最终的购买决定，是由十几轮对话中逐步显现的情感信任、认知调整、偏好重塑共同促成的。在这个过程中，哪些部分该归因给“平台的推荐算法”，哪些部分该归因给“用户的自主判断”，哪些部分该归因给“对话情境中非特定地归属于任何一方的涌现效应”——这些问题没有客观的答案。你可以设计一种归因算法，把70%的权重分配给AI的推荐话术、30%分配给用户的自我驱动。但这只是算法设计者的一种选择，不是对客观因果链的忠实记录。换一个设计者，换一种归因逻辑，比例可以是50-50。这就是记录权在绵延交互中失效的根本原因：记录权要求存在一条可以被客观记录的因果链，但在绵延交互中，因果链的可记录性本身瓦解了——不是因为技术不够好，记录不到，而是因为在那里根本不存在一个“唯一正确的因果链”等待被记录。因果链不是被“掩盖”了，而是被“消解”了。

在这个困境中，一种新的权力形态被实际的经济竞争逼迫着发生出来。在2022年到2024年间，主流电商平台上的品牌方——尤其是那些高度依赖线上渠道的消费品牌——开始发现一件让他们不安的事：他们投在传统广告位上的预算，转化效果在系统性地下滑。用户不再从广告位点击进入商品页，而是直接打开AI导购对话。在对话中，AI可能推荐了A品牌，也可能推荐了B品牌——而无论是A还是B品牌方，都无法从后台数据中获知自己的品牌为什么被推荐或未被推荐。他们从后台看到的只是：来自AI导购渠道的订单在涨，来自传统搜索广告渠道的订单在跌，而他们无法通过调整广告出价来影响前者。（这一现象的综合重构基于多个电商平台公开发布的行业趋势报告与第三方数字营销服务商的分析，不取自任何单一品牌的内部运营数据，且经匿名化处理。）在这个压力下，一部分品牌方被迫转向了一种新策略：不再争夺“广告位”，而是争夺“被AI在对话中自然地提及”。他们开始投入资源去优化自己产品在AI训练数据中的“在场质量”——确保自己的产品描述、用户评价、售后数据足够

丰富和正向，使得模型在训练过程中“自然地”将他们的品牌与某些高频情境关联起来。他们无法购买这种关联，他们只能通过持续优化自己的数据在场来“养”出这种关联。

这就是“在场权”发生的具体机制。它不是被任何一个人发明出来的新概念，而是被品牌方在归因失效的逼迫下，被迫摸索出的一种替代性权力逻辑。在场权不依赖于事后追溯因果链，而是靠事前占据价值生成所必须的现场条件。一个AI导购之所以能截走原本属于某品牌的客流，不是因为它抢了谁的归因链，而是因为它在用户还在犹豫不决的那个微妙时刻，用一段精准的对话把用户的信任接过去了。在那段对话发生的当下，品牌、传统电商平台、比价网站——它们都不在场。在场的只有用户和AI。价值就在这一在场性中被捕获。

为了不让这个概念悬浮在修辞层面，需要为它划定严格的属性边界。在场权区别于记录权的关键维度如下：记录权的他者是广告主——它的权力运作是通过日志（物件）进行事后追溯；在场权的他者是那些不在场的竞争者——它的权力运作是通过参数空间（痕迹）进行事前潜入。记录权的失效条件是“归因链断裂”，而在场权的失效条件则是“在场被商品化”——即当平台允许品牌直接付费购买模型内部的倾向性，从而把“在场”重新交易化成另一种形式的“记录”时，在场权就不再成其为在场权了，它会退化为记录权的一种更隐蔽的变体。这一自我限定，是把在场权从一个敏锐的观察变成一个可被精确使用的分析工具的必要步骤。

5.3 竞争战场的转移：从屏幕空间到参数空间

记录权向在场权的转移，同步地带来了竞争战场的转移。互联网经济时代，竞争的核心战场在屏幕空间。首页排名、搜索结果置顶、应用商店推荐位——这些都是竞争者们拼死争夺的领土。领土的逻辑是清晰的——一个位置只能被一个主体占据。竞争是可见的，竞争结果是可测量的。AI经济时代，竞争的核心战场从可见的屏幕空间，转移到了不可见的高维参数空间。在互联网经济中，如果你想要让更多用户看到你的商品，你的主要手段是去争夺曝光位。在AI经济中，“曝光”本身的定义都变得模糊了。一个用户可能从未在传统意义上“看到”你的商品——他只是在和AI对话的过程中，AI在第三轮推荐中自然地带出了你的品牌。这个推荐是怎么来的？它不是因为你购买了某个关键词广告位，而是因为你的品牌在模型的高维参数空间中，与“操作简单”“适合老年人”“售后服务好”这些用户此刻最关心的特征维度形成了高度关联。这种关联，是模型在训练过程中从海量数据中自行建立起来的。它没有办法被简单地购买——它只能被“训练进去”。竞争变成了在不可见的参数空间中争夺那一套痕迹所赋予的倾向性。这个位置不是一件可以买卖的广告位商品，它是被数据、训练策略、模型架构共同塑造出的一种隐性倾向。竞争的结果变得更难以测量——因为“赢”不再是“我的广告被展示了X次”，而是“在无数次对话中，我的品牌被AI以‘最自然的推荐’的方式带出来的次数”，而这个数字是没有人能准确统计的。

六、主体的重塑：从决策劳动到信任委托

6.1 广袤中的决策劳动：一种被经济构造原理维持的官能

互联网经济赋予经济主体一个极具辨识性的处境：被抛入一个广袤但透明的选择空间。说“广袤”，是说选项在数量上几乎是无限的——电商平台上数百万SKU，任何关键词都能搜出成千上万条结果。说“透明”，则是在互联网经济的规则体系下，每一个选项的属性至少在原则上是可以被获知和比较的：价格可以一键排序，销量与好评率公开展示，商品参数以结构化的表格呈现。这种信息环境迫使经济主体练就了一套内在程序——把一堆杂乱的选项放在一起，从中识别出关键差异、排除噪声、做出有根据之权衡。这套程序不是天生的，也不是一旦学会就终身不变的；它是在广袤信息平原中被日常经济行为持续操练、持续维持的。每一次搜索、每一次比价、每一次差评细读，都是一次微小的决策劳动——它让人累，但它也让人持续地确认自己作为决策者的身份。

6.2 肥沃中的信任委托：另一种官能会发生

AI经济给予经济主体的处境，从“广袤”翻转到了“肥沃”。“肥沃”指的是：选项不再排山倒海地摊在你面前，让你自己去比较、去挣扎。每次你看到的，只是那一个、或很少几个被精心筛选出来的、与你的历史行为与当下情境高度匹配的选项。一个对话式AI导购不会给你列出五十款血压计让你自己去比参数——它直接告诉你“根据你父母的情况，我推荐这款”，然后附上三条简短的理由。决策的费力程度被降到了前所未有的低位：你不需要比价，不需要读上百条评论，不需要在不同参数之间做艰难的权衡。最好的选项被端到你面前，你所需要做的只是一个动作——接受，或者拒绝再换一个。

当经济主体持续地生活在这样一个“被喂答案”的环境中时，一种什么样的内在变化正在发生？当那种需要在多个选项中挣扎、比较、取舍的日常经济行为大规模退场时，那种由这些行为所操练的决策官能，就从经济生活中逐渐闲置了。这不是“退化”——退化预设了一种规范性的“应有状态”。更准确地说，这是一种生态学的替代：在一个选项爆炸的信息平原上，不具备发达决策官能的个体无法有效行动，因此这个环境持续选择并维持着这种官能；在一个由AI精心筛选推荐的环境中，决策官能的使用频率和必要性大幅降低，而另一种官能——信任委托的能力——的回报率急剧上升。能够快速识别AI推荐的可靠性、能够有效地用自然语言表达自己的模糊需求、能够在系统出现错误时迅速察觉并切换到人工干预——这些能力，取代了传统的比价、筛选、权衡，成为新经济环境中的适应优势。

必须明确指出：这一主体的重塑，并非独立于“规则崩解”之外的另一个社会学后果。它正是规则崩解在“人”的层面的具体发生。决策劳动的存在条件是什么？是一个广袤但透明的选择空间——在那个空间里，选项是物件（可以被清点、被比较、被排序），规则也是物件（价格排序规则、好评率计算规则），主体在这些物件的支撑下，把自己的理性形态塑造成了“在物件之间做出有根据的权衡者”。当规则崩解为痕迹，交易崩解为绵延场，那个做出“有根据之权衡”的物件基础就解除了。主体不再站在一堆可以被清点的选项面前，而是站在一段绵延的交互面前——在那段交互里，选项不是被呈现的，而是被那一瞬间“养”出来的。主体无法再做有根据的权衡，因为需要被“据”的那些物

件，已经不成其为物件了。信任委托，正是在这个物件真空中被逼出来的替代性官能。它不是主体的主动选择，而是主体在物件化理性形态失去操作对象后，被新环境逼着长出来的一种适应形态。

七、回应严肃的反驳

7.1 反驳一：人类制度从来都是“黏性”的——这一翻转是否被夸大了？

一个根本性的反驳直指本文核心论点的适用范围：人类法官的司法判决、管理者的商业直觉，都同样无法被还原为清晰的规则条文。判决中有法感、有直觉、有对案件整体情境的不可言明的把握——这些难道不也是痕迹式的？如果人类制度从来都在“物件的笼子里装着痕迹式的内容”，那么本文把“物件vs痕迹”这对概念赋予互联网经济与AI经济之间，是否只是把一个原本存在于一切人类制度中的张力，夸大成了两个经济时代之间的断裂？

本文对这一反驳的回应是：人类法官的痕迹式判断，与AI模型的痕迹式决策，确实同样位于“决策无法被逐条还原为清晰物件”这一端。但有一笔关键差别，正是这笔差别构成了断裂的硬核：人类法官的痕迹，被套在一个物件的制度笼子里——这个笼子不要求笼子里的东西透明，但它要求笼子里的东西说出一个可供共同体争辩、复审、推翻的“说理”。法律体系用一个制度性的物件外壳——必须给出书面说理、说理必须经得起上级审查——把法官内心无法被物件化的痕迹，强行收束在了一个可以被公共讨论、被制度性挑战的框架内。这个笼子的构成要素，不仅是“必须说理”这一程序性要求，更关键的是说理内容的可争辩性：败诉方可以基于判决书中写下的理由提起上诉，上级法院可以基于这些理由来审查原判决是否“适用法律错误”或“违反逻辑法则”。说理不仅是说出来了，而且是拿到了一个可以被他人用公开标准检验的场域里。

AI模型目前没有这个笼子。一个深度学习推荐模型决定不向某商家分配流量时，它不需要给出任何“说理”。即使监管部门或商家要求解释，模型能给出的也只是事后用简单模型拟合出来的、与模型实际决策逻辑无关甚至矛盾的“解释”。这个解释不能被争辩——因为不存在一个可以依据这些解释来“推翻模型决策”的制度程序。不能被复审——因为模型可以在下一次训练中悄悄改变权重，没有一条清晰的“判例”可以被援引。因此，“物件vs痕迹”的对立，不是在说“以前的东西都是物件，未来的东西都是痕迹”——从来都不存在纯粹透明的物件规则，任何人类制度中都填塞着无法被完全写定的痕迹。这对概念切开的裂缝在于：互联网经济用一套制度性的“物件外壳”把痕迹收束在了一个可以被公共讨论和制度性挑战的框架之内；AI经济中，这个物件外壳在技术上被消解了，而替代它的制度笼子尚未被发明。断裂不在痕迹本身的有无，而在“笼子”的有无。

7.2 反驳二：这只是技术不成熟——等解释技术足够好，痕迹就能被翻译回物件

第二个值得认真对待的反驳是：可解释AI是一个活跃的研究领域，监管机构也在推行算法透明与算法审计制度。未来很可能出现一种融合形态：AI系统能够提供足够高质量的事后解释，以至于它在功能上重新实现了互联网经济的可审计性。本文在第二章已经给出了对这一反驳的技术回应——事后解释工具与模型实际决策逻辑之间的系统性偏离，不是暂时现象，而是结构性的。这里要补充更根本的一点：即使解释技术取得了突破，也不意味着痕迹被“还原”回了物件。因为解释是事后附着上去的，不是从痕迹内部提取出来的。它不改变痕迹本身的存在方式。一个被贴了说明牌的痕迹，仍然是痕迹——它的存在方式不曾因此改变。这意味着，即使监管机构能够在功能上重新建立一套“可审计性”，这套可审计性的本体论基础已经变了。它不再是“因为规则是物件所以可审计”，而是“我们

在痕迹旁边立了一套制度性的物件外壳，让它可以被当作物件来对待”。后者的稳定性，取决于那套制度外壳本身的持续维护——一旦维护松懈，痕迹就会裸露出来，物件外壳的脆弱性就会暴露。

一个有力的经验佐证来自金融监管实践。在模型风险管理（Model Risk Management, MRM）的正式要求下，金融机构的信用评分模型必须接受监管审查，审查的核心之一就是要求开发者提供每一个输入特征的影响方向与理论依据。这是典型的物件式审计要求——监管者想要看到“收入水平对信用评分的影响方向是正向的”这样一条可以被单独写定、单独验证的规则陈述。工程师为了满足这一要求，使用SHAP值等工具生成特征重要性报告。但工程师心里清楚——有经验的模型开发者知道——SHAP值报告的不是模型“内部遵循的规则”，而是对模型行为在某个局部的线性近似。换一种特征排序方法，换一组参照样本，同一个特征的“重要性”可以发生显著漂移。监管者拿到了一份看起来像“规则物件”的文件，可以把它装进审计档案；但工程师知道，文件所描述的那个“规则”，在模型内部并不存在。这就是在痕迹旁边立起的物件外壳——它有实践价值，它让审计得以运转，但它的脆弱性在于：一旦监管进一步深化，要求基于这些特征重要性报告来“校正”模型中的某些“错误规则”，校正的效果将不可预测。那层物件外壳在承担“推演责任”时会暴露它下面的痕迹本质——你不能基于SHAP值去“修改”模型，就像你不能基于沙滩上的指示牌去“修改”脚印的走向。

7.3 反驳三：断裂还是加速？——互联网经济早已是痕迹式的

第三个反驳来自经济学经验研究。真实的大型互联网平台，早在引入深度学习之前，其规则体系就已经具备了本文归于AI经济的许多特征——归因链的任意性、边界的内生模糊性、交易记录的建构性。如果这些“痕迹化”的表征在互联网经济时代就已经大量存在，那么本文把“断裂”定位在互联网经济与AI经济之间，是否把一场已经在发生的渐变叙事，硬切成了戏剧性的断裂？

本文的回应不是否认这些表征的存在，而是指出它们与本文诊断的“规则崩解”之间存在一个决定性的差别。在互联网经济中，归因链虽然武断，但它仍然有一条可以被争辩的制度性边界——广告主可以抗议“最后点击归因”，要求改用“多触点归因”，平台必须给出一种归因逻辑，且这种逻辑在商业合同中是可以被约定、被审计的。审核虽然从来不是纯粹二值的，但审核团队仍然依据一份被写定的审核指南做出判断，且被误判的用户可以通过申诉程序要求人工复审。这些表征的共性在于：它们虽然暴露了物件范式的极限，但它们仍然被物件式的制度工具所收束。互联网经济时代的“痕迹化”是局部性的、过程性的、被制度笼子约束着的。AI经济的断裂在于：深度学习使得这个收束的物件式框架本身，在技术上失去了立足点。你无法要求模型给出“归因逻辑”——因为没有一条可以被写定的归因逻辑存在在那里。互联网经济的痕迹化是“笼子里的痕迹”，AI经济的痕迹化是“笼子本身也在消融”。断裂不在痕迹的有无，而在是否还存在一个可以收束痕迹的制度笼子。

这一点在另一个正在进行的经验拉扯中被进一步照亮。在某些高度受监管的金融科技领域，监管机构正在强行用一套物件式的审计要求去套痕迹式的AI模型。MRM审查要求模型开发者提供每一个输入特征的理论依据和影响方向说明，这实际上是在模型周围强行建起一个物件式的制度笼子。这场拉扯——监管者要求物件，工程师知道那边只有痕迹——恰好证明了本文的判断不是在宣称“所有AI经济都已变成痕迹”，而是在揭示一种结构性的张力：AI经济的技术趋向是痕迹化，而制度体系正在反向拉扯，试图强行维持物件外壳。拉扯的现场，就是本文所追问的那个发生学现场：痕迹化的技术压力与物件化的制度惯性之间，正在积累出什么样的新矛盾？制度笼子在这次拉扯中，会被加固到足以重

新收束痕迹，还是会在拉扯中暴露出自身的历史条件性，从而在某个临界点上被放弃？这些追问不是本文能够回答的，但它们比“我们需要新笼子”这种治理呼吁，更接近发生学该做的事。

八、结论：规则崩解之后，追问不休

本文试图论证：互联网经济与AI经济之间的断裂，不是“规则从透明变为黑箱”、不是“从可解释变为不可解释”——这些描述都还站在“规则是物件”的河岸上，它们只在问“规则长什么样”。本文要追到河的源头去问：规则是怎么被做出来的？对这个问题的回答揭示了一次本体论翻转。在互联网经济中，规则是物件——被制定、被书写、被修改、被审计。在AI经济中，规则崩解为痕迹——它从未被“制定”过，它只是模型在数据上走过之后留在参数空间里的那一套残留。这套残留有效用，但它没有独立的身份，无法被单独指向、单独操作、单独审计。“修改规则”这个动作在物理上失去了操作对象——不是因为规则藏得太深，而是因为那个可以被捏住、被抽取、被改写的独立物件，在源头上就没有被生产出来。

这一翻转不是进步也不是退步——它是互联网经济自身矛盾推动下的发生过程。当规则多到不再能被单独操作时，物件范式的实质就已经萎缩了；AI把这一萎缩推进到了它的逻辑终点——连物件的外表都不再维持。规则从未以物件的形态存在过，它们从一开始就是痕迹。

走到这里，有一种诱惑是把结论收束为对“新治理工具”的呼吁——在痕迹的外部建立一套新的制度笼子。这个方向诚然是实践者会选择的第一条路。但本文必须追问一个比“如何建新笼子”更根本的问题：旧笼子的崩解，并不自动意味着新笼子的建成；它意味着连“需要笼子”这个需求本身，都要被重新问题化。那个被我们称为“治理者”的主体——它的专业知识、它的管辖权边界、它与被治理者之间的可辨识关系——这些都是在物件范式所支撑的制度构造中被发生出来的。当物件范式崩解时，被拖入追问的不仅是“治理对象是什么”，而且是“治理者是谁”以及“治理这个动作是什么意思”。在痕迹式的世界里，“治理”还能作为一个可以被单独辨认、单独授权的社会动作存在吗？还是它将与它所试图治理的那些痕迹一起，融进一片整体漂移的参数空间里——在其中，治理与被治理的边界本身也变成了一套痕迹？这个问题没有现成的答案，但它是发生学追问在痕迹化世界里必须抵达的位置。

本文的核心论断可以被一条清晰的证伪条件所检验：如果未来十年内，出现了一种AI经济系统，它能够把每一次定价、每一次推荐、每一次交互背后实际起作用的决策逻辑——不是事后的近似拟合、而是实际运作的机制——用一套可以被独立调取、独立修改、独立审计的规则物件完整地实现出来，并且这种物件化丝毫不削弱该系统的经济效率，那么本文的核心判断就是错的。因为那意味着痕迹可以被无损翻译回物件，“规则崩解”并没有真正发生——痕迹只是物件被暂时藏起来的假象。而如果这两者真的能无损互译，那么AI经济就只是更快、更复杂的互联网经济，而不是另一套构造经济世界的方式。本文的论证倾向于相信这种无损互译是不可能的——不是因为技术还不够好，而是因为物件与痕迹属于不同的存在类别。物件是被制造的，痕迹是被留下的。你能在痕迹旁边建一座制度性的物件笼子，但你不能把痕迹变成物件。正如你能在海滩的脚印旁插一块指示牌，但你不能把脚印变成石板路。未来需要的不是在痕迹表面铺更多石板，而是在没有石板的世界里，试探出不必依赖石板的行走方式——同时保持对“行走”这个动作本身的追问：在没有石板的世界里，“行走”是什么意思？

注释

- [1] 材料说明：本文所涉谷歌RankBrain案例（2.3节）与电商平台AI导购案例（3.2节）均非基于一手实地调研，而是综合自公开技术文档、媒体报道与行业共识的综合性重构，并经简化和匿名化处理，作为思想拓展性质的例示引用。其中，AI导购案例中的具体对话与“信任变化”判断是基于对话逻辑的叙事重构，不声称其基于情感计算或经验实验的精度。以上案例旨在帮助呈现论证逻辑，不反映任何特定企业的内部数据或实际运行细节。

参考文献

- Aneesh, A. (2009). “Global Labor: Algocratic Modes of Organization.” *Sociological Theory*, 27(4), 347–370.
- Google. (2016). *How Search Works*. Google.
- Lessig, L. (1999). *Code and Other Laws of Cyberspace*. Basic Books.
- Lipton, Z. C. (2018). “The Mythos of Model Interpretability.” *Communications of the ACM*, 61(10), 36–43.
- Ribeiro, M. T., Singh, S., & Guestrin, C. (2016). “Why Should I Trust You?” Explaining the Predictions of Any Classifier. In *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining* (pp. 1135–1144). ACM.